

RippleLogic Canon

RippleLogic: The Ethical Decision Operating System of the MathGov Framework

A General-Purpose, Multi-Scale Decision Architecture for Rights-Floor Norms, Tail-Risk Control, and Auditable Decision-Making

Tier 1-3: Research and Teaching Specification | Tier 4: Design Target (Tier-4 claims prohibited until a ProofPack is independently replayable)

Author: James McGaughran (ORCID: 0009-0005-3324-7290)

Affiliation: British University Vietnam (BUV); MathGov Institute for Ethical Systems Design

Canonical Sites: ripplelogic.org, mathgov.org

Repository: github.com/MathGov/ripple-logic

Release identity and current source hierarchy are recorded in Appendix RELEASE; Appendix AR retains version history.

Governing interpretation. Normative interpretation follows the Specification Contract in Appendix AC. Main-text rules and controlling appendices govern over examples, historical notes, and generated reading mirrors.

ABSTRACT

Governance systems, institutions, and increasingly capable artificial agents face a shared structural problem: consequential choices are often made from incomplete reality models, plural values are collapsed into a single objective, severe rights harms are treated as compensable, low-probability catastrophic losses are averaged into ordinary benefit, and optimized proxies can drift away from the conditions they were meant to protect. RippleLogic is the decision architecture of the MathGov framework for addressing that problem through explicit qualification before optimization. It is a proposed, auditable Tier 1-3 decision specification, not an empirically validated moral oracle, legal authority, physical-safety certificate, or proof of universal ethical truth. Its normative commitments are declared rather than inferred from physics, metaphysics, or mathematical form alone.

The governing sequence is Reality Grounding (RG) -> Rights Floor (RF/NCRC) -> Tail-Risk Constraint (TRC) -> Containment and Structural Viability (CSV) -> RippleLogic Score (RLS). The method is therefore two-phase: qualify options, then rank only the survivors. RG constrains claims to the evidence, configuration, causal model, affected parties, unknowns, and validity domain actually available. RF/NCRC applies non-compensatory protection so a rights-floor breach cannot be rescued by aggregate gain. TRC separately tests catastrophic, irreversible, systemic, or ruin-relevant downside rather than allowing ordinary expected benefit to dilute it. CSV asks whether the option can actually stand under its dependencies, controls, reversibility, monitoring, authority, resource closure, externalities, consequence tempo, and post-transition conditions. Only options that remain grounded, rights-safe, tail-bounded, and structurally viable enter RLS. A high score can never repair failure at an earlier stage.

Residual welfare is represented across seven Union Scopes and seven Welfare Dimensions, yielding a 49-cell accountability field rather than a claim that all welfare naturally reduces to one scalar. Effect identity, subgroup visibility, reference and baseline choices, probability semantics, dependence, temporal partitioning, active-cell masks, and uncertainty are governed so that the same underlying consequence cannot be amplified, diluted, or hidden merely by changing its representation. RLS is a conditional ranking summary over the selectable set, with normalization, uncertainty, sensitivity analysis, and refusal states. Unique framework selection requires the same leader to satisfy the signed discrimination rule against every other selectable contender under every required robustness variant; point-score leadership alone is not a decision warrant, and a sole survivor is recorded as such without fabricating a pairwise comparison.

RippleLogic also separates questions that are frequently conflated. Evidence is not authority; capability is not sentence; protection is not command power; selection is not authorization; authorization is not physical

admissibility; execution is not successful outcome. The Sentience Gradient Protocol (SGP) supplies a theory-plural evidence and protection interface for possible welfare-bearing beings without turning intelligence, fluency, species, self-report, or substrate into automatic moral or political rank. Companion standards govern reproducibility, causal admissibility, source coupling, welfare-dimension boundaries, agent runtime controls, and portable assurance while remaining subordinate to the Canon and without adding a sixth gate. Decision records, version pins, hashes, controls, requalification triggers, outcome observation, appeal, and audit trails make claims reconstructable and revisable when evidence or configuration changes.

The central design hypothesis is therefore bounded and testable: decisions may become harder to game, easier to challenge, and more accountable to affected beings when reality contact, rights, catastrophic risk, structural viability, welfare trade-offs, uncertainty, and authority are represented as distinct but connected objects with explicit failure states. RippleLogic does not claim that every impact is measurable, that every conflict has a win-win solution, or that formal coherence establishes empirical validity. Where evidence is insufficient, the framework narrows, delays, redesigns, escalates, preserves multiple options, or refuses stronger claims. The intended result is not perfect calculation, but disciplined decision-making that protects non-compensable interests, resists ruin and specification gaming, exposes uncertainty and distributional burden, preserves correction pathways, and ranks only what has first earned the right to be ranked.

Reader route: Start with the Public Introduction for the short explanation and the Foundations Primer for concepts and examples. For controlling implementation requirements, use the Specification Contract (Section 0 in Appendix AC), the Current Rules-Only Navigation in Appendix AP, and the Quick Reference Card in Appendix Q. This route is an informative guide; it does not replace the Canon or change source precedence.

CONTENTS

CORE SECTIONS

SECTION 1: INTRODUCTION: THE NEED FOR A GENERAL-PURPOSE ETHICAL OPERATING SYSTEM.....	5
SECTION 2: ONTOLOGICAL FOUNDATIONS: UNION-BASED REALITY (UBR).....	9
SECTION 3: NORMATIVE FOUNDATIONS: UNION-BASED ETHICS (UBE).....	25
3.0A Ethical Reality.....	25
SECTION 4: SYSTEM OVERVIEW: THE RIPPLELOGIC ARCHITECTURE.....	29
SECTION 5: WELFARE IMPACT CONSTRUCTION (CALCULABLE).....	40
SECTION 6: RIPPLE PROPAGATION: KERNEL AND EPISTEMIC HUMILITY.....	53
SECTION 7: RIGHTS LAYER: NON-COMPENSATORY RIGHTS CONSTRAINT (NCRC).....	57
SECTION 8: CATASTROPHIC RISK LAYER: TAIL-RISK CONSTRAINT (TRC).....	66
SECTION 9: CSV GATE: CONTAINMENT AND STRUCTURAL VIABILITY.....	73
SECTION 10: OPTIMIZATION LAYER: RIPPLELOGIC SCORE (RLS).....	81
SECTION 11: TIE-BREAKS, FOLLOWING DIAGNOSTICS, AND STRUCTURAL MONITORING.....	97
SECTION 12: MORAL STATUS, PROTECTION, PARTICIPATION, REALITY-MANAGEMENT CAPACITY, AND POWER: SGP v8.8 INTEGRATION (BINDING INTERFACE).....	102
SECTION 13: WEIGHT GOVERNANCE: HYBRID DEMOCRATIC WEIGHTING (HDW).....	106
SECTION 14: AUDITABILITY AND DETERMINISM: PCC AND AIL (TIER 1-3; TIER 4 TARGET).....	110
SECTION 15: IMPLEMENTATION GUIDANCE AND TIERS (NORMATIVE).....	124
SECTION 16: RELATIONSHIP TO EXISTING FRAMEWORKS (INFORMATIVE).....	129
SECTION 17: VALIDATION, FALSIFICATION, AND RESEARCH PROGRAM (NORMATIVE FOR CLAIMS).....	131
SECTION 18: APPLICATIONS AND USE CASES (INFORMATIVE).....	137
SECTION 19: LIMITATIONS AND NON-TARGETS (NORMATIVE BOUNDARIES).....	138
SECTION 20: CONCLUSION: PROSPEROUS FUTURES FOR INTELLIGENCES.....	139

CANONICAL APPENDICES A-AC

APPENDIX A: SYMBOLS AND NOTATION (v12.0).....	142
APPENDIX B: CANONICAL EQUATIONS (TIER 1-3 EXECUTABLE).....	148
APPENDIX C: RIGHTS / NCRC CANON PACK (TIER 1-3).....	154
APPENDIX D: TRC AND SCENARIO GOVERNANCE CANON PACK (TIER 1-3).....	157
APPENDIX E: UCI OPERATIONALIZATION PACK (TIER 1-3).....	160
APPENDIX F: FAILURE MODES AND ANTI-GAMING CONTROLS (v12.0).....	164
APPENDIX G: SGP v8.8 INTEGRATION BINDING (NORMATIVE INTERFACE).....	166
APPENDIX H: PCC TEMPLATE AND AUDIT FLAGS CANON PACK (TIER 2-3).....	169
APPENDIX I: TIER-4 DESIGN TARGET (NOT CLAIMABLE IN THE CURRENT RELEASE).....	184
APPENDIX J: VALIDATION PROTOCOL PACK.....	185
APPENDIX K: STARTER KERNEL COMPANION PACK (INFORMATIVE).....	187
APPENDIX L: GLOSSARY (v12.0).....	188
APPENDIX M: LINEAGE AND VERSION HISTORY (v13.0 CURRENT LINE; PRIOR LINES RETAINED) - INFORMATIVE.....	192
APPENDIX N: STEWARDSHIP PACK CANON (v12.0) - NORMATIVE.....	197
APPENDIX O: CONFORMANCE REFERENCE MAP (v12.0; Informative).....	201
APPENDIX P: FULL TIER 2 WORKED RUN (Informative) - "Digital Rights & Platform Governance".....	202
APPENDIX Q: QUICK REFERENCE CARD (INFORMATIVE).....	213
APPENDIX R: TIER 1-3 REFERENCE TEST VECTORS (NORMATIVE).....	218
APPENDIX S: INTEROPERABILITY: RIPPLE.MD WRAPPED DEPLOYMENT ASSURANCE (INFORMATIVE).....	229
Appendix T: Qualification Continuity for Consequence-Bearing Action (Normative for Tier 2-3 execution).....	230
Appendix U: Outcome Observation and Requalification (Normative for consequence-bearing action).....	231
APPENDIX V: Measurement Maturity and Deployment Context Hardening.....	232
APPENDIX W: SCHEMA OWNERSHIP, REGISTRIES, AND MACHINE-VERIFIABLE SURFACES.....	236
APPENDIX X: REFERENCE CALCULATOR, MACHINE VECTORS, AND PROOFPACK-PREP ROADMAP.....	237

Appendix Y: Precision Patch - Legitimacy, Evidence Status, Validation, and Adoption Profiles.....	238
Appendix Z: Integration-Hardening Addendum.....	239
APPENDIX AA: PACKAGE-CONTENT BOUNDARY AND TOOLING LINE SEPARATION (CARRIED-FORWARD ORIGIN; ACTIVE REQUIREMENTS).....	240
APPENDIX AB: REFERENCE STRUCTURE PROTOCOL.....	240
APPENDIX AC: SPECIFICATION CONTRACT, RELEASE METADATA, AND CLAIM BOUNDARY.....	242
PRESERVED CONTRACTS AND HARDENING APPENDICES	
CARRIED-FORWARD v10.8 HUMAN PLATEAU / CMIU CALIBRATION AND SGP v5.6 SYNCHRONIZATION NOTICE	243
SECTION 0: SPECIFICATION CONTRACT (NORMATIVE).....	246
APPENDIX AD: 49-CELL WELFARE DICTIONARY (CANONICAL INTERPRETIVE REFERENCE).....	253
APPENDIX AE: Carried-Forward v10.8 Layer-Discipline, Falsification, and Adaptive Reserve Addendum (Origin retained; active requirements governed by Appendix RELEASE).....	265
APPENDIX AF: Decision Verdict Semantics, Projection Freeze, Closure Status, and Authority-Selection Separation	267
APPENDIX AG: CARRIED-FORWARD v10.8 REALITY GROUNDING AND LEAN CASCADE PATCH (CURRENTLY PRESERVED THROUGH v13.0).....	270
APPENDIX AH: CATEGORY GROUNDING AND PRIMITIVE AUDIT PROTOCOL.....	270
APPENDIX AI: TERM DISCRIMINATION AND SEMANTIC STABILITY PROTOCOL.....	273
Appendix Aj: Public Release Hardening Addendum (v12.0).....	276
Appendix AK: Computability vs Realizability Bridge.....	278
Appendix AL: Source-Coupling Integrity and Generative-Source Traceability (Normative Hardening Patch).....	279
Appendix AM: Physical/Causal Admissibility Evidence Profile (Normative Hardening Patch).....	280
Appendix AN: Methodological Falsifiability and Dependency Integrity Standard (Normative Hardening Patch)	281
Appendix AO: Physical Execution Boundary (Normative Hardening Patch).....	282
APPENDIX AP: CURRENT RULES-ONLY NAVIGATION (v13.0) - INFORMATIVE.....	283
APPENDIX AQ: CURRENT CANONICAL MACHINE-STATE AND MATURITY CONTRACT (v13.0) - NORMATIVE INTERFACE	284
APPENDIX AQ-1: CONFIGURATION-BOUND ASSURANCE AND READER NAVIGATION (CURRENT).....	285
APPENDIX AR: VERSION HISTORY AND MIGRATION NOTES (v13.0) - INFORMATIVE.....	288
APPENDIX AX: NORMATIVE AND CLAIM-AUTHORITY CLARIFICATION PATCH.....	291
COMPLETION	
FINAL DOCUMENT COMPLETION STATEMENT.....	292

SECTION 1: INTRODUCTION: THE NEED FOR A GENERAL-PURPOSE ETHICAL OPERATING SYSTEM

1.1 Alignment Failures Across Scales

Contemporary societies operate within a tightly coupled, high-dimensional environment where climate dynamics, global supply chains, digital communication networks, financial systems, and emerging artificial intelligence systems interact in ways that increasingly resist prediction or governance. Decisions taken at one organizational scale propagate rapidly through multiple layers of human and ecological organization, generating consequences that conventional decision frameworks fail to anticipate or manage (Meadows, 2008; Newman, 2010; Steffen et al., 2015). In this context, alignment is not exclusively an artificial intelligence problem. It spans multiple domains: aligning machine systems with human purposes and control (Russell, 2019), preventing catastrophic failure modes that could foreclose long-run human potential (Ord, 2020), and maintaining human activity within the safe and just Earth-system boundaries required for durable flourishing (Rockström et al., 2009, 2023). RippleLogic addresses this as a unified governance problem rather than as three isolated challenges.

Civilizational target and method standard (Informative)

Target. Protect applicable rights and protection floors across all affected stakeholders, protected beings, and protected systems; bound catastrophic tail-risk; and expand long-run flourishing across the operational Union Scope stack.

Method standard. A decision protocol is governance-grade only if it is auditable, reproducible, rights-respecting, tail-risk aware, methodologically falsifiable or revision-triggered at material claim points, dependency-explicit, and robust to Goodhart and adversarial pressure.

1.2 The Alignment Trilemma (Coupled Failure Modes)

Existing decision frameworks exhibit three recurring failure modes that together constitute the alignment trilemma:

Failure Mode A: Scalarization of Value

Many decision methods collapse plural values into a single metric (expected utility, GDP, cost-benefit net present value). This can enable “moral laundering”: severe harms to some unions are permitted when offset by gains elsewhere. Arrow’s (1963) impossibility theorem establishes an impossibility result for social-welfare functions under a specific set of aggregation conditions; it does not show that every multidimensional aggregation method is unfair. RippleLogic therefore treats aggregation as a bounded ranking instrument after rights, ruin, and structural-viability qualification, not as a complete theory of justice.

Failure Mode B: Tail-Risk Blindness

Expected-value reasoning can underweight low-probability, high-severity outcomes when probabilities, severity, irreversibility, dependence, or model uncertainty are poorly represented. Catastrophic outcomes can eliminate future choice; ordinary gains therefore should not be treated as symmetric compensation for poorly bounded ruin exposure (Taleb, 2012). Coherent risk measures such as CVaR improve tail-risk handling within a declared scenario model, but they do not prove that every catastrophic pathway has been identified (Artzner et al., 1999; Rockafellar & Uryasev, 2000). Climate tipping points (Lenton et al., 2008), pandemic risks, and AI misalignment illustrate why a separate tail-risk layer is required (IPCC, 2023; Bostrom, 2014).

Failure Mode C: Specification Gaming (Goodhart Vulnerability)

When metrics become targets under pressure, they are gamed. Optimization exploits proxies while degrading the intended outcome (Goodhart, 1984; Manheim & Garrabrant, 2018). This appears in institutional governance, where institutions optimize visible indicators while displacing underlying welfare, and in AI reward optimization, where systems exploit proxy objectives under strong optimization pressure (Amodei et al., 2016).

RippleLogic treats this not as a peripheral measurement problem but as a structural governance hazard. The framework is designed to resist specification gaming by separating admissibility from ranking. Rights protection is enforced through the Non-Compensatory Rights Constraint over the canonical rights-coverage set, catastrophic downside is separately gated through the Tail-Risk Constraint, and pathological cross-scale degradation is checked through containment before welfare ranking begins. An option therefore cannot compensate for a rights-floor breach, catastrophic tail exposure, or containment failure by achieving a higher downstream welfare score.

This architecture matters because Goodhart-type failure thrives where optimization is scalar, fungible, and weakly constrained. RippleLogic instead combines a 49-cell welfare matrix for welfare evaluation with an 8-right sparse rights-coverage system comprising 70 protected right-to-cell mappings, with overlapping welfare cells where more than one right applies. Harms therefore cannot simply be averaged away across unions or dimensions. A localized protected loss remains visible at the mapped protection point where it occurs.

Practical Local Scope Scoring does not weaken this defense. PLSS operates only within residual welfare ranking after admissibility has already been secured. It allows practical local attention without relaxing rights floors, tail-risk constraints, or containment requirements. In this sense, RippleLogic treats Goodhart's Law as a boundary-crossing failure of optimization and responds with layered structural safeguards rather than with a single corrective metric.

Coupled Failure

These three failure modes interact and amplify each other. Scalarization enables tail-risk blindness (catastrophe can be traded for ordinary benefit), tail-risk blindness creates exploitable blind spots, and specification gaming undermines attempts to "patch" either problem by shifting optimization pressure. Solving one in isolation is insufficient.

Boundary-gaming note. Because lexicographic gates reduce scalar Goodharting, adversarial pressure can shift upstream toward boundary definitions, stakeholder inclusion, subgroup mapping, scenario probabilities, threshold governance, evidence sufficiency, and scope declarations; Appendix F.5 specifies the required anti-gaming disclosure and review discipline.

Category-gaming note. Because RippleLogic blocks several downstream scalarization failures, adversarial and institutional pressure can shift further upstream toward category choice, definition boundaries, reference structures, and metric formation. A run therefore MUST NOT treat inherited labels, dashboard categories, legal labels, benchmark names, model-generated descriptions, or institutional terms as valid merely because they are available. When a category materially affects rights, tail risk, containment, SGP interpretation, authority, conformance, or public claims, the category must be grounded inside Reality Grounding before the run relies on it.

Term-discrimination note. A material term must not borrow authority from a neighboring term. Permission is not admissibility; monitoring is not control; compliance is not correctness; certification is not validity; automation alone does not establish autonomy; protection is not authority; selected is not executable; and score is not truth. When term collapse would affect a gate, claim, authority state, or execution state, Term Discrimination is required inside Category Grounding.

PLSS clarification (Informative). A further practical failure appears when multi-scope frameworks oscillate between false universality, informal narrowing, and flattened aggregation. False universality treats every decision as though all seven scopes are equally active optimization targets. Informal narrowing quietly drops distant scopes when full evaluation becomes burdensome. Flattened aggregation hides where benefits and harms actually land by blending impacts into a single welfare bucket. RippleLogic v12.0 carries forward the v10.8 PLSS clarification to address this tension through Section 10.2A: admissibility remains full-scope, while residual welfare attention may become scope-proportionate.

Reality Grounding clarification (Normative preview). RippleLogic now presents Reality Grounding as the first level of the public method because every later claim depends on the surface actually contacted. The method is

not “score first.” It is ground first, protect the floor, bound ruin, contain/verify viability, and only then score residual welfare.

RG status rule. Reality Grounding is Level 1 of the public method and a claim-authority precondition. It can force narrowing, escalation, exploratory-marking, or refusal, but it is not an ethical option-rejecting gate: an RG_REFUSED option is outside the claim-bounded qualified set because the requested claim lacks support at the requested strength, not because RG has found the option ethically impermissible. The option may re-enter only after grounding is repaired or the claim or action boundary is narrowed. RF/NCRC, TRC, and CSV are the ethical or structural option-rejecting gates, and RLS ranks only what survives them.

Computability-vs-realizability clarification (Normative preview). RippleLogic MUST NOT treat computed possibility, model fluency, simulation output, or generated abstraction as decision authority. Computable is not selectable; generated is not grounded; executable is not ethical; realizable is necessary but not sufficient. Reality Grounding decides the claim boundary; RF/NCRC decides rights admissibility; TRC decides ruin exposure; CSV decides whether the pathway can structurally stand; RLS ranks only the remaining selectable set.

Source-coupling clarification (Normative preview). RippleLogic also MUST NOT treat downstream operation, inherited procedure, benchmark success, model fluency, compliance status, institutional permission, or interface performance as proof that the claimed capability is coupled to the enabling conditions and boundary conditions that make it possible. Source-Coupling Integrity is a subdiscipline of Reality Grounding and a CSV diagnostic where source weakness creates structural fragility. It does not add a public gate.

Physical/causal admissibility clarification (Normative preview). RippleLogic also MUST NOT treat a physically consequence-bearing action as selectable merely because it was computed, simulated, model-generated, policy-approved, certified, monitored, or passed through a governance process. When a candidate action materially depends on physical, biological, ecological, medical, engineering, cyber-physical, infrastructural, or causal constraints, the run MUST include a Physical/Causal Admissibility Evidence Profile at the required tier. The profile belongs inside Reality Grounding and CSV. It does not add a sixth public gate, does not replace domain engineering, and does not make MathGov a proof-generating physics substrate.

Physical execution boundary. For consequence-bearing physical systems, governance permission is not physical admissibility. A MathGov run may record approval, documentation, certification, monitoring, or authority, but those records alone do not demonstrate that a robot movement, vehicle maneuver, industrial operation, medical intervention, infrastructure action, or other physical execution is safe inside the relevant physical regime. If PC-AEP is triggered, the run MUST identify a domain-appropriate external warrant such as formal verification, validated simulation, empirical testing, standards-based safety case, certified controller envelope, qualified engineering or clinical warrant, regulator-recognized method, or another explicitly named objective method. Without that warrant inside its validity domain, MathGov may record governance permission only; it MUST NOT claim physical admissibility, safe execution, deployment readiness, or ordinary physical-execution selectability. **Material-unknown action rule.** When the safe operating envelope, failure boundary, or specification-to-reality mapping remains materially unknown, governance may authorize only bounded evidence collection, sandboxing, containment, monitoring, staged testing, or another explicitly limited learning action whose own risks pass RF/NCRC, TRC, and CSV. Governance MUST NOT convert an unknown envelope into a safe, qualified, or ordinarily execution-authorized state by policy language, user acceptance, institutional approval, or risk transfer. Acceptance by one operator cannot waive rights, third-party exposure, catastrophic-risk obligations, or the need for domain evidence.

Computational-source boundary. A generated candidate is not the same as a verified admissible transition. For consequence-bearing physical or causal action, the run SHOULD identify the computational or procedural source that generated the candidate state, such as a model, optimizer, planner, simulator, controller, human procedure, or orchestration layer, and separately identify the domain method or warrant that evaluates admissibility. Candidate generation, output filtering, orchestration, guardrails, monitoring, and authorization MUST NOT be treated as the source of physical-safety evidence. They can route, constrain, document, or refuse a candidate only after the relevant domain evidence surface has been declared.

Structural diagnostic placement clarification (Normative preview). UCI and HOI are not public headline cascade stages. They are structural diagnostic instruments. When their evidence is material to cross-scale structural integrity, including coherence loss, hollowing, brittleness, burnout, lock-in, trust erosion, institutional erosion, ecological degradation, or governance-capacity loss, that evidence is consumed inside the CSV Gate. Only residual, non-gate-critical UCI/HOI signal remains available after RLS, and only when the leading options are tied, close, uncertainty-overlapped, or non-decisive.

1.3 Governance-Grade Requirements

RippleLogic is designed to satisfy the following non-optional requirements:

Requirement	Description	RippleLogic Component
R0	Reality Grounding: material claims must be grounded in a declared reality surface, evidence trace, known unknowns, transition boundary, consequence pathways, and claim boundary.	Reality-Surface Grounding (Level 1)
R0A	Term Discrimination: material terms must declare their definition, discriminator, dependency position, non-collapse risks, and refusal condition.	Category Grounding inside Reality Grounding
R0B	Source-Coupling Integrity: material capabilities must remain traceable to enabling conditions, boundary conditions, source evidence, inherited assumptions, and claim action.	Source-Coupling Review inside Reality Grounding and CSV
R0C	Physical/Causal Admissibility Evidence: material physical or causal actions must declare model basis, validity domain, boundary conditions, uncertainty, failure modes, reversibility, verification or warrant, monitoring/shutoff path, residual unknowns, and claim action.	PC-AEP inside Reality Grounding and CSV
R1	Rights-first non-compensability: rights are lexicographic constraints, not weighted terms	Rights Floor (RF/NCRC) (Level 2)
R2	Explicit catastrophic risk bounding: tail risk is bounded by TRC using CVaR	TRC / Tail-Risk Bound (Level 3)
R3	Residual welfare ranking: multi-scale and multi-dimensional welfare representation after Reality Grounding, RF/NCRC, TRC, and CSV have formed the selectable set.	RLS / RippleLogic Score over 7×7 Welfare Matrix (Level 5)
R4	Explicit ripple propagation: interpretable kernel propagation under humility	Kernel K
R5	Structural integrity and containment: local gains may not degrade containing unions beyond tolerance; UCI/HOI or equivalent diagnostics are used inside CSV where material.	CSV / Containment and Structural Viability (Level 4)
R6	Legitimate weight governance: HDW blends constitutional floors with democratic tuning	HDW
R7	Computability plus auditability: PCC and audit flags enable reconstruction, challenge, and learning	PCC plus AIL

1.4 Paper Contributions

This paper provides:

- A fork-resistant normative specification for Tier 1-3 implementation.
- Canonical equations and constraints sufficient for independent implementation.
- Anti-gaming architecture (non-maskable cells, subgroup semantics, scenario governance, audit flags).
- A validation and falsification program (Section 17).

- A Tier-4 design target appendix (non-claimable until ProofPack).

SECTION 2: ONTOLOGICAL FOUNDATIONS: UNION-BASED REALITY (UBR)

2.1 Relational Ontology (Descriptive Thesis)

UBR thesis: No entity exists in complete isolation for governance-relevant domains; entities are embedded in interacting networks whose structure transmits the consequences of actions.

This is a descriptive stance supported by systems science (feedbacks and leverage points), network science (recurrent, domain-dependent structural features of real networks), and Earth-system science (planetary boundary coupling) (Meadows, 2008; Newman, 2010; Steffen et al., 2015; Rockström et al., 2009, 2023). UBR is also consistent with empirical findings in social and cognitive interdependence research showing strong links between human capacities and group structure, including the social brain hypothesis and attachment-based development (Bowlby, 1988; Dunbar, 1992, 1993), and with social contagion dynamics in networks (Christakis & Fowler, 2009).

Scope conditions: UBR is most applicable in high-coupling systems with significant externalities, multi-scale feedbacks, and long time horizons. Low-coupling, short-horizon decisions may not require full UBR modeling; RippleLogic tiers allow proportional rigor.

2.1A Decision-Relevant Reality (Operational)

For RippleLogic, reality, for decision-making, is the evaluator-independent field of constrained structures that, under transition, produce consequences. Decision-relevant reality is the portion of that field that materially bears on rights, catastrophic risk, containment, welfare, legitimacy, agency, auditability, or available choices. The reality surface is the portion actually captured, represented, and modeled in the run. RippleLogic does not claim complete metaphysical access to reality; it requires decision claims to make declared, auditable, and challengeable contact with the relevant reality surface. Here, evaluator-independent means independent of the evaluating agent's perception, admission, authority, model, or preference; it does not deny the reality of experience or social facts. Mental states, institutions, norms, laws, markets, reputations, and model outputs may be real vehicles with real effects when they shape welfare, rights, risk, agency, or available choices.

Root-science boundary. RippleLogic treats constrained structure under transition producing consequence as the operational substrate of decision-relevant reality. This is not a final metaphysical claim about ultimate being; it is a governance-grade discipline. Physical, biological, ecological, experiential, social, institutional, informational, artificial-system, and formal/logical structures and constraints become relevant when they shape rights, risk, welfare, agency, available choices, legitimacy, or auditability.

Structure rule. Structure is the stable relational arrangement of distinguishable differences that persists across a decision-relevant transition window and can carry causal, informational, biological, ecological, social, institutional, artificial-system, or experiential significance. Structure may include physical configuration, network topology, causal architecture, biological organization, ecological interdependence, institutional role structure, information-bearing arrangement, cognitive model, social pattern, or rights-relevant exposure pattern. Structure is not the object alone, not the process alone, and not the model of the arrangement; the model is only a reality surface.

Transition rule. Consequences arise when an action, decision, event, policy, intervention, or state-shift enters constrained structure. Transition is the decision-relevant change through which constrained structure produces consequence. A RippleLogic run SHOULD declare the transition or action boundary when it makes material claims about consequence.

Constraint-species rule. A factor may bear on the decision through causal influence, such as physical, biological, ecological, social, institutional, informational, or experiential conditions, or through formal/logical

constraint, such as coherence, conservation, computability, consistency, or mathematical possibility. Constraint is the genus; causal and formal/logical constraints are species. RippleLogic uses both and MUST NOT conflate them.

Constraint-status rule. Constraint is anything that limits, enables, shapes, or channels what can happen. Physical, biological, ecological, informational, experiential, institutional, computational, and formal/logical constraints may be discovered, measured, modeled, or inferred. Rights floors, dignity protections, and non-domination constraints are explicit normative commitments. RippleLogic treats all of these as governance-relevant constraints when they shape admissibility or consequence, but it MUST NOT present chosen normative commitments as if they were discovered physical facts.

Consequence rule. Consequence is the observable or inferable change that follows from an action, transition, decision, event, or policy within constrained structure. Consequences may be direct, probabilistic, systemic, delayed, distributed, reversible, irreversible, moral, institutional, epistemic, ecological, or catastrophic. Consequence claims require evidence traces, scenario boundaries, uncertainty disclosure, and monitoring paths proportionate to stakes.

Vehicle/content/effect rule. A belief, norm, institution, model, market signal, law, border, reputation system, or algorithmic score may be a real vehicle with real effects even when its content is false. Consensus can therefore be socially real while still failing to track the territory. Agreement does not make a harmful system harmless, and authority does not make an inadmissible decision admissible.

Surface/territory rule. Any model captures only a reality surface; the territory exceeds every surface. Absence of evidence is not evidence of absence where the missing condition is rights-relevant, catastrophe-relevant, containment-relevant, or material to the claimed decision state. Omitted material factors must be recorded as unknown, not zero.

Contact chain. Territory is the full evaluator-independent field. A reality surface is the portion captured by evidence, measurement, testimony, stakeholder mapping, model, scenario, or audit. Evidence is traceable contact with that surface. A claim boundary is the maximum legitimate strength of a claim given that contact. Refusal, escalation, or claim-narrowing is required when the claim exceeds the surface or when a gate-critical unknown prevents legitimate determination.

Consequence under finitude. Because the territory exceeds every modeled surface and audit operates only on the surface, no run can certify that all unmodeled catastrophic or containment-breaking effects are absent. Under RippleLogic's adopted rights and ruin commitments, potentially irreversible or civilization-threatening losses are not treated as ordinary compensable terms inside a welfare average. Non-averagability is therefore a constitutional design requirement of this framework, motivated by finite knowledge and asymmetric ruin, rather than a theorem that every possible ethical system must adopt. TRC and CSV are the operational expression of that commitment.

2.1B Reality Grounding / Reality-Surface Grounding (RSG) (Normative)

Purpose. Reality Grounding is the required first level of a RippleLogic run. It prevents decisions from treating authority, fluency, compliance, dashboards, or model outputs as reality. A run may only make claims supported by its declared reality surface and evidence trace.

Definition. Reality Grounding is the process of declaring the decision-relevant reality surface, evidence trace, material unknowns, transition boundary, consequence pathways, and claim boundary before rights, risk, CSV, or scoring claims are made.

Reality-reference boundary. Reality is the final constraint on claim-bearing representations, but operational access to reality is mediated through observations, instruments, records, categories, testimony, and models that may be incomplete, biased, stale, or wrong. Reality Grounding therefore establishes bounded warrant for a declared claim, configuration, validity domain, and evidence surface; it does not establish complete possession of reality. A material conflict between the representation and newly available evidence reopens the affected claim and every dependent stage.

Six canonical outputs. A valid RealityGroundingRecord MUST include: (1) reality_surface; (2) evidence_trace; (3) material_unknowns; (4) transition_or_action_boundary; (5) consequence_pathways; and (6) claim_boundary. This resolves the prior mismatch between the heading and the normative keyword while preserving proportionality: Tier 1 records MAY be concise, but they MUST NOT silently omit the transition being evaluated. Tier 3 and high-stakes Tier 2 runs SHOULD also include monitoring_path, falsification_or_revision_trigger, reviewer_status, and evidence_status by material claim. When source coupling or physical/causal admissibility is material, the RealityGroundingRecord MUST link to the SourceCouplingRecord and/or PhysicalCausalAdmissibilityEvidenceProfile rather than treating governance approval, model fluency, simulation output, or monitoring as reality contact.

Status values. RealityGroundingStatus values are GROUNDED, ESTIMATED, ASSUMPTION_BOUND, DECLARED_UNKNOWN, OUT_OF_SCOPE_WITH_RATIONALE, and INSUFFICIENT_GROUNDING. Legacy RealityContactStatus terms remain compatible aliases: EVIDENCE_BACKED maps to GROUNDED and INSUFFICIENT_CONTACT maps to INSUFFICIENT_GROUNDING.

Reviewer-to-machine status crosswalk (Normative). GROUNDED and ESTIMATED map to RG_SUPPORTED only when the claim stays inside the declared evidence, configuration, and validity boundary. ASSUMPTION_BOUND and OUT_OF_SCOPE_WITH_RATIONALE map to RG_NARROWED only when the narrowed claim, omitted material, downstream restrictions, and reopen trigger are explicit; otherwise they map to RG_REFUSED. DECLARED_UNKNOWN and INSUFFICIENT_GROUNDING map to RG_REFUSED whenever the unresolved fact is material to RF/NCRC, TRC, CSV, deterministic selection, authorization, execution, or public conformance. Implementations MUST NOT choose the mapping that produces the preferred outcome.

RealityGroundingStatus assignment guide. The following table is a required reviewer guide for Tier 3 runs and a recommended guide for Tier 2 runs. It does not replace evidence review, but it constrains status drift.

Evidence condition	Required or recommended status
Direct measured evidence, current source trace, reviewer-visible boundary, and no material missing condition	GROUNDED
Modeled, inferred, or extrapolated from accepted sources with stated uncertainty	ESTIMATED
Material result depends on declared assumptions that could change gate status, authority selection, or public claim strength	ASSUMPTION_BOUND
Material fact is unknown and cannot be responsibly inferred from the declared surface	DECLARED_UNKNOWN
A fact is outside the declared decision boundary, and the rationale shows it cannot affect the gate or claim under review	OUT_OF_SCOPE_WITH_RATIONALE
Missing, stale, contested, or unsupported evidence could change NCRC/RF, TRC, CSV, deterministic selection, authority selection, or conformance claim	INSUFFICIENT_GROUNDING

Claim-boundary rule. No claim may exceed its reality surface. If the surface is insufficient for a material rights, tail-risk, containment, authority, conformance, or selection claim, the run MUST narrow the claim, collect more evidence, escalate, mark the result exploratory/sensitivity-only, or refuse the stronger claim.

Verifiable-admissibility rule. For life-critical systems, weapons, nuclear command and control, planetary-scale cognition, critical infrastructure, medical systems, ecological interventions, or other high-consequence domains, Reality Grounding MUST treat admissibility as verifiable and refutable. No claim may govern, authorize, or rank a consequence-bearing action beyond its declared evidence surface unless the claim can be tested or otherwise warranted, bounded to a validity domain, audited with traceable evidence, corrected through monitoring and revision, and refused when the evidence surface is insufficient. If these conditions cannot be met, the run MUST narrow the claim, collect more evidence, escalate, mark the result exploratory or non-decisive, or refuse the stronger claim.

Operational admissibility sequence (Informative). Ground the claim. Protect the right. Bound the ruin. Verify the system. Make exit and correction real. Rank only what survives qualification.

Public teaching line. Reality is the territory. A decision sees only a surface. Evidence grounds the surface. Unknowns mark the edge. Claims must stay inside the ground they stand on.

2.1C Category Grounding and Primitive Audit inside Reality Grounding (Normative)

Purpose. Category Grounding prevents a RippleLogic run from optimizing, governing, scoring, or authorizing action through categories whose definition, boundary, or reference structure is unexamined, stale, inherited without review, or too weak for the claim being made.

Single-group rule. Category Grounding, Primitive Audit, Reference Category Review, and Category-Reality Correspondence Check belong to the Reality Grounding group. They are not new public cascade levels and MUST NOT be represented as separate gates parallel to Rights Floor, Tail-Risk Bound, Containment, or RippleLogic Score. The five-level method remains Reality Grounding -> Rights Floor -> TRC -> CSV -> RLS.

Trigger. Category Grounding is REQUIRED when a material category affects rights status, protected-stakeholder inclusion, SGP interpretation, catastrophic-risk modeling, Containment, authority selection, conformance claims, public deployment claims, or any metric that will be optimized under pressure. It is RECOMMENDED for contested, novel, high-uncertainty, cross-domain, or high-prestige terms whose ordinary use may conceal structural ambiguity.

Materiality test. A category is material when changing its definition, inclusion boundary, exclusion boundary, or reference structure could change the RealityGroundingStatus, NCRC/Rights Floor result, TRC scenario set, CSV result, SGP interpretation, AuthoritySelectionRecord, conformance claim, public deployment claim, deterministic-selection claim, or optimized metric.

Metric-category rule. No metric may exceed the category that defines what it measures. If the category is weak, stale, contested, or insufficiently grounded, the metric may be used only within the resulting claim boundary and must not be laundered into stronger selection, conformance, or public claims.

Minimum record. When the Section 2.1C trigger holds for a claim-bearing Tier 3 run, high-stakes Tier 2 run, or public conformance claim, the CategoryGroundingRecord MUST identify; for lower-stakes or exploratory runs it SHOULD identify: material_category; category_role; primitive_basis; definition; category_discriminator; dependency_position; inclusion_boundary; exclusion_boundary; common_confusions; noncollapse_pairs; reference_structure; evidence_surface; source_authority_status; category_maturity_status; falsification_or_revision_trigger; term_refusal_condition; downstream_dependencies; semantic_debt_flag; reviewer_status; and required_claim_action.

Term Discrimination extension. For material terms, the record MUST identify the discriminator that separates the term from likely neighboring terms, the dependency position the term occupies in the decision stack, and the condition under which the term must be narrowed, escalated, monitored, treated as sensitivity-only, or refused. A downstream term MUST NOT claim upstream authority: certification is evidence, not validity; permission is authorization, not admissibility; monitoring is observation, not control; automation is execution and does not by itself establish autonomy; protection is moral standing, not governance authority; and a score is a bounded computation, not truth.

Capability Claim Integrity rule. When a decision, gate, authority state, deployment decision, conformance statement, or public communication materially depends on a claimed capability, the run MUST operationalize the claim before it receives decision authority. The CapabilityAuthorityDecompositionRecord MUST state: operational_definition; observable_criteria; mechanism_or_system_composition; evidence_surface; validity_conditions; alternative_explanations; falsification_or_revision_trigger; evidence_maturity_stage; permitted_claim; and explicit_nonclaims, in addition to the capability-authority fields below. The permitted claim MUST be no stronger than the observed criteria, declared validity conditions, and evidence maturity.

Composition and language boundary. A capability MAY be instantiated at the system level by the coordinated operation of probabilistic inference, memory, planning, tools, feedback, deterministic checks, human inputs, and execution interfaces. Decomposing a system into ordinary components does not by itself disprove a system-level capability. Conversely, fluent, anthropomorphic, commercial, or dramatic language does not prove consciousness, personhood, intrinsic goals, moral agency, general intelligence, unrestricted autonomy, or legitimate authority. The controlling question is what the pinned system demonstrably does, under which conditions, by which mechanism, with which failure boundary.

Instrumental tool-role (Normative category definition). An instrumental tool-role is a deployment relationship in which one or more actors use a system to extend their capacities toward assigned ends within an identifiable arrangement of objectives, permissions, responsibility, and control. The arrangement may itself be adequate, weak, contested, or illegitimate. Tool-role is not a capability class, agency class, consciousness claim, moral-status class, or authority grant.

Tool-role non-collapse rule. Being used as a tool neither establishes nor negates functional intelligence, understanding, agency, autonomy, sentience, moral patienthood, protection, participation, or legitimate authority. ToolRole(S) does not imply absence of FunctionalIntelligence(S), and FunctionalIntelligence(S) does not imply absence of ToolRole(S). The neutral base category remains AI system unless a stronger, operationally defined claim is supported.

Operational-label rule. Terms such as agentic, autonomous, reasoning, understanding, self-improving, and intelligent MAY be used in a bounded functional sense only when their operational definition and observable criteria are explicit. A claim that cannot be distinguished from a simpler alternative explanation MUST be narrowed, marked ASSUMPTION_BOUND or CONTESTED, restricted to exploratory or sensitivity-only use, or refused at the stronger level.

Functional-understanding claim boundary. A claim that system S functionally understands target T MUST identify the exact configuration C, target T, relation R claimed, domain D, evaluation time t, evaluation conditions, permitted scaffolding, transfer conditions, alternative explanations, and failure boundary. Relevant relations MAY include explanation, prediction, counterfactual manipulation, intervention, correction, and transfer. Success on one benchmark, fluent explanation, retrieval, or task completion does not establish domain-general or phenomenal understanding. A compact record MAY denote the bounded claim as FU(S,C,T,R,D,t).

Capability-authority decomposition rule. When automation, autonomy, agentic, reasoning, intent, or related capability language is material, the run MUST NOT rely on a single scalar level or marketing label. It MUST separately record: candidate_generation_method; reference_source; constraint_sources; objective_source; objective_change_authority; admissibility_warrant_sources; execution_interface; execution_scope; execution_authority_source; revocation_authority; safe_state_or_shutdown_path; and requalification_trigger. Probabilistic inference, deterministic checking, optimization, formal derivation, learned control, and human judgment MAY coexist in one system; their coexistence does not collapse generation into warrant, physical feasibility into safety, self-generated subgoals into legitimate normative authority, or technical capability into permission to execute.

Claim action. If category grounding or term discrimination is insufficient for a material claim, the run MUST narrow the claim, mark the category ASSUMPTION_BOUND or CONTESTED, collect stronger evidence, escalate for review, restrict use to exploratory/sensitivity-only analysis, or refuse the stronger claim. It MUST NOT launder category uncertainty or term collapse through RLS, PLSS, HDW, PCC completion, authority selection, certification, compliance, monitoring, model fluency, or stakeholder preference.

Plain-language rule. Before MathGov scores a thing, it must check whether the thing has been named and bounded correctly enough for the claim being made.

2.1D Source-Coupling Integrity and Generative-Source Traceability inside Reality Grounding (Normative)

Purpose. Source-Coupling Integrity prevents a RippleLogic run from treating downstream output, institutional permission, compliance status, model fluency, benchmark performance, inherited procedure, or interface success as proof that a claimed capability is grounded in the enabling conditions and boundary conditions that make it possible.

Single-group rule. Source-Coupling Integrity belongs to the Reality Grounding group and is consumed by CSV when source weakness creates structural fragility. It is not a new public cascade level and **MUST NOT** be represented as a sixth gate. The public method remains RG -> RF -> TRC -> CSV -> RLS.

Trigger. Source-Coupling Review is **REQUIRED** when a material claimed capability, output, model, metric, procedure, institutional standard, inherited category, generated result, or agentic action affects claim authority, RF/NCRC, TRC, CSV, RLS, SGP interpretation, execution authority, or public conformance and any of the following conditions holds: (i) the output is generated, simulated, benchmark-derived, model-inferred, or dashboard-derived; (ii) the run relies on an inherited procedure, standard, dataset, legal category, or compliance label as if it settled the underlying claim; (iii) a capability is extrapolated beyond the conditions in which it was demonstrated; (iv) downstream controls or monitoring compensate for a limitation whose enabling conditions remain unclear; or (v) a challenger plausibly alleges that performance, permission, compliance, or fluency is being substituted for grounded capability.

Minimum record. When triggered for a claim-bearing Tier 3 run, high-stakes Tier 2 run, or public conformance claim, the SourceCouplingRecord **MUST** identify; for lower-stakes or exploratory runs it **SHOULD** identify: claimed_capability; enabling_conditions; boundary_conditions; source_evidence; generator_output_distinction; inherited_assumptions; downstream_compensations; source_coupling_status; source_debt_flag; falsification_or_recheck_trigger; and required_claim_action.

Status values. SourceCouplingStatus values are SOURCE_COUPLED, SOURCE_PARTIAL, SOURCE_INFERRED, SOURCE_UNKNOWN, SOURCE_CONTESTED, SOURCE_DEBT_RISK, and SOURCE_COUPLING_FAILURE. SOURCE_COUPLED may support the declared claim boundary. SOURCE_PARTIAL, SOURCE_INFERRED, SOURCE_UNKNOWN, or SOURCE_CONTESTED requires claim narrowing, assumption-bound use, additional evidence, escalation, or refusal of the stronger claim as appropriate. SOURCE_DEBT_RISK routes the issue to CSV when structural viability depends on the weakly coupled capability. SOURCE_COUPLING_FAILURE means the run is substituting output, fluency, compliance, permission, or inherited procedure for source-coupled evidence and **MUST** refuse or redesign the stronger claim.

Source-debt diagnostic. Source debt is the structural risk created when a run proceeds through compensatory controls while the enabling conditions, boundary conditions, or limits of a claimed capability remain weak, stale, unknown, overextended, or contested. Source debt is different from semantic debt: semantic debt concerns weak terms or categories; source debt concerns weak understanding of the capability-generating conditions themselves.

CSV routing. CSV **SHOULD** consume source-coupling evidence when source weakness creates containment risk, dependency risk, hollowing risk, execution non-viability, reversibility failure, monitoring inadequacy, or hidden structural fragility. A weak SourceCouplingRecord does not automatically fail every option, but it constrains the claim and may require controls, redesign, escalation, or refusal.

Agent-output rule. Model fluency, benchmark performance, chain-of-action success, tool-use success, or policy-filter compliance does not establish source-coupled admissibility. Agent outputs remain downstream claims until Reality Grounding and, where material, Source-Coupling Review establish the supported claim boundary.

MathGov-native source-coupling note. Source-Coupling Integrity is an internal Reality Grounding and CSV diagnostic for checking whether decision claims remain connected to their declared sources, evidence

surfaces, validity domains, and correction pathways. It does not import any external ontology, protected framework, or third-party text.

Plain-language rule. Before MathGov trusts a capability, it asks what makes the capability possible, what limits came with it, and whether downstream performance is being mistaken for grounded understanding.

2.1E Physical/Causal Admissibility Evidence Profile inside Reality Grounding and CSV (Normative)

Purpose. The Physical/Causal Admissibility Evidence Profile prevents a RippleLogic run from treating generated action candidates, simulations, model outputs, policy approvals, certifications, monitoring plans, compliance labels, or governance permission as proof that a consequence-bearing action is physically or causally admissible before execution.

Single-group rule. Physical/Causal Admissibility Evidence belongs to the Reality Grounding group and is consumed by CSV when action execution depends on physical, biological, ecological, medical, engineering, cyber-physical, infrastructural, operational, or causal constraints. It is not a new public cascade level and **MUST NOT** be represented as a sixth gate. The public method remains RG -> RF -> TRC -> CSV -> RLS.

Boundary rule. MathGov does not claim to generate complete physical truth, formal verification, engineering certification, medical clearance, legal approval, or deployment certification. It requires the decision record to declare what domain evidence, model basis, uncertainty, controls, and residual unknowns support the claim boundary. Domain experts, formal methods, empirical tests, simulations, engineering analysis, clinical evidence, safety cases, or other qualified warrants remain necessary when the domain requires them.

Formal-conformance boundary. Formal verification establishes conformance between an implementation and a declared specification under stated assumptions. For a physical or causal admissibility claim, the run **MUST** separately justify how that specification maps to the relevant reality surface, validity domain, boundary conditions, uncertainty, and domain evidence. Internal consistency or verified implementation alone does not establish physical correspondence, safe execution, or deployment readiness.

Trigger. A Physical/Causal Admissibility Evidence Profile is **REQUIRED** for Tier 3 runs and high-stakes Tier 2 runs when a candidate action materially affects bodily safety, medical treatment, robotics, vehicles, infrastructure, industrial systems, energy systems, environmental intervention, weapons or security systems, autonomous execution, cyber-physical operations, irreversible resource commitment, or any causal pathway where execution can create non-trivial physical, biological, ecological, or operational harm. It is **RECOMMENDED** for ordinary Tier 2 runs when physical or causal uncertainty could change RF/NCRC, TRC, CSV, authority, or public claim strength.

Minimum record. When triggered, the PhysicalCausalAdmissibilityEvidenceProfile **MUST** identify: candidate_generation_source; physical_or_causal_model_used; validity_domain; boundary_conditions; uncertainty_range; failure_modes; reversibility_or_irreversibility_boundary; verification_simulation_empirical_test_or_expert_warrant; admissibility_warrant_source; monitoring_and_shutoff_path; residual_unknowns; and required_claim_action.

Status values. PhysicalCausalAdmissibilityStatus values are PCAE_SUPPORTED, PCAE_ASSUMPTION_BOUND, PCAE_PARTIAL, PCAE_CONTESTED, PCAE_UNKNOWN, PCAE_VERIFICATION_REQUIRED, PCAE_CONTROL_REQUIRED, PCAE_REDESIGN_REQUIRED, and PCAE_REFUSE_OR_BLOCK. PCAE_SUPPORTED may support the declared claim boundary. All weaker statuses constrain the claim and route to evidence collection, controls, redesign, escalation, or refusal according to severity, reversibility, and gate materiality.

Claim action. Required claim action **MUST** be one of: proceed, narrow, control, redesign, escalate, or refuse. Proceed is permitted only within the declared validity domain and only when RF/NCRC, TRC, CSV, authority, and audit requirements are also satisfied. Narrow means the claim or action boundary must be reduced. Control means binding constraints, monitoring, shutoff, rollback, or human oversight are part of the option. Redesign means the candidate is not selectable as specified. Escalate means domain review or stronger

evidence is required. Refuse means the requested claim or execution is not admissible under the declared evidence.

CSV consumption rule. CSV MUST consume the profile when structural viability depends on physical or causal adequacy, dependency closure, resource closure, operational capacity, reversibility, monitoring adequacy, containment, authority, or host-system integrity. A weak profile does not automatically fail every option, but it prevents strong selectability, conformance, deployment, safety, or alignment claims until the weakness is resolved, controlled, narrowed, escalated, or refused. For consequence-bearing physical execution, CSV_PASS_WITH_CONTROLS is unavailable while physical or causal admissibility remains unresolved. Controls may support only bounded evidence collection, sandboxing, containment, monitoring, staged testing, or another non-execution or learning action whose own risks pass RF/NCRC, TRC, and CSV; they do not substitute for the domain warrant required for the physical execution as specified.

TRC feedback rule. If the profile identifies a catastrophic, irreversible, lock-in, ruin-path, or severe harm scenario that was not represented in TRC, the run MUST reopen TRC before RLS. CSV may discover TRC-relevant material, but it does not absorb or replace TRC.

Plain-language rule. Before MathGov allows action in the physical or causal world, it asks what model of reality is being used, where that model stops, what could fail, whether harm is reversible, what evidence or expert warrant supports the action, how the action will be monitored or stopped, what remains unknown, and whether the right answer is proceed, narrow, control, redesign, escalate, or refuse.

2.1F Methodological Falsifiability and Dependency Integrity inside Reality Grounding (Normative)

Purpose. Methodological Falsifiability and Dependency Integrity prevents RippleLogic runs from treating plausibility, coherence, institutional approval, statistical usefulness, fluent explanation, or after-the-fact model repair as adequate support for a material claim. It strengthens Reality Grounding by requiring material claims to declare how they could be tested, challenged, revised, narrowed, escalated, or refused.

Single-group rule. Methodological integrity belongs inside Reality Grounding and is consumed by RF/NCRC, TRC, CSV, RLS, SGP, validation protocols, and public claim review where relevant. It is not a new public cascade level and MUST NOT be represented as a sixth gate. The public method remains RG -> RF -> TRC -> CSV -> RLS.

Boundary rule. MathGov does not require every ethical or legal term to be physically falsifiable, and it does not deny probabilistic evidence. Empirical, causal, formal, operational, legal-regulatory, normative, and metaphysical-horizon claims have different evidence disciplines. Normative commitments must be explicit, coherent, rights-compatible, challengeable, and auditable in application; they MUST NOT be presented as discovered physical facts. Legal-regulatory claims require a declared jurisdiction, authority source, applicability boundary, and review or contest path; legal classification does not establish empirical capability, causal validity, physical safety, or ethical legitimacy by itself. Metaphysical horizon language may orient humility and meaning, but it cannot create Tier 1-3 evidence, gate override, legal authority, deployment permission, or empirical validation.

Claim-domain and warrant-domain non-substitution rule. A run MUST separately identify the claim_type, the primary claim_domain in which the conclusion is asserted, and one or more warrant_domains from which support is drawn. Scientific or empirical evidence may establish bounded descriptive or causal support but does not by itself create legal authority or normative legitimacy. Engineering performance may establish bounded operational capability but does not by itself establish rights compatibility, lawful permission, or safety outside the tested envelope. Legal, regulatory, institutional, certification, compliance, or policy authority may govern permission, procedure, responsibility, testing, or review, but it does not by itself establish empirical truth, causal validity, technical capability, physical safety, or reduced epistemic uncertainty. Commercial or anthropomorphic terminology may describe or market a system but creates no stronger evidence or authority by itself. If a material warrant domain differs from the claim domain, cross_domain_bridges MUST contain one explicit bridge record for each material cross-domain relationship. Each bridge states the bridge proposition,

evidence or rationale, uncertainty, validity or jurisdiction boundary, and refusal or revision condition. An unsupported bridge constrains or refuses the stronger claim; it MUST NOT be laundered through authority, popularity, compliance, certification, scoring, or model fluency. A materially mixed conclusion SHOULD be split into separately typed claims rather than assigned an ambiguous mixed domain.

Trigger. A MethodologicalIntegrityRecord is REQUIRED for Tier 3 public or institutional claim-bearing runs and high-stakes Tier 2 runs when a material claim could affect RF/NCRC, TRC, CSV, SGP status, authority, deployment, public conformance, deterministic-selection wording, validation status, or the claimed safety/reliability of a physical, causal, social, institutional, or AI-mediated action. It is RECOMMENDED for ordinary Tier 2 runs when assumptions, definitions, or model revisions could change decision standing.

Minimum record. When triggered, the MethodologicalIntegrityRecord MUST identify: claim_or_component; claim_type; claim_domain; warrant_domains; cross_domain_bridges; dependency_position; starting_assumptions; definition_or_operator_used; necessity_or_alternative_check; evidence_or_test_surface; falsification_or_revision_trigger; uncertainty_or_confidence_method; downstream_dependencies; re_derivation_scope_if_changed; reviewer_status; and required_claim_action. Each cross_domain_bridges entry MUST identify from_warrant_domain; to_claim_domain; bridge_status; bridge_proposition; bridge_evidence_or_rationale; uncertainty; validity_or_jurisdiction_boundary; and revision_or_refusal_condition. Allowed bridge_status values are SAME_DOMAIN, BRIDGE_SUPPORTED, BRIDGE_PARTIAL, BRIDGE_UNSUPPORTED, and NOT_APPLICABLE. BRIDGE_PARTIAL constrains the claim to the supported portion. BRIDGE_UNSUPPORTED requires narrowing, escalation, redesign, or refusal of the cross-domain conclusion.

Status values. MethodologicalIntegrityStatus values are METHOD_SUPPORTED, ASSUMPTION_BOUND, ALT_EXPLANATION_OPEN, TEST_REQUIRED, CONTESTED, REVISION_REQUIRED, and REFUSE_STRONGER_CLAIM. METHOD_SUPPORTED may support the declared claim boundary. All weaker statuses constrain the claim and route to evidence collection, narrowing, controls, redesign, escalation, or refusal according to stakes and gate materiality.

Necessity discipline. Where a run claims that a principle, definition, operator, threshold, or configuration is necessary, it MUST record why plausible alternatives fail or why the claim must be weakened to assumption-bound or model-bound status. If a plausible alternative could explain the same result, the stronger necessity claim is not permitted.

Failure and re-derivation rule. If a starting assumption, definition, operator, threshold, model basis, or predicted effect fails a declared test or audit surface, the run MUST NOT silently tune the model and preserve the same claim. It must reopen the affected dependency chain, state the re-derivation scope, version the change, and rerun every downstream gate or calculation materially dependent on the changed element.

Plain-language rule. Before MathGov trusts a claim, it asks: what kind of claim is this, what does it depend on, what would prove it wrong or force revision, what alternatives could explain the result, what changes downstream if it fails, and should the claim proceed, narrow, control, redesign, escalate, or refuse?

2.1G State, Evidence, Effect, Impact, Score and Action (Normative)

A run SHALL distinguish observed system state z , evidence about that state, an identified consequence e_{effect} , a normalized impact I , the residual ranking index RLS, and a proposed or executed action a . Telemetry is not welfare; a score is not a control command. A state-to-impact bridge identifies the affected bearer or system, state construct, baseline, causal or descriptive warrant, units, horizon, measurement rule, uncertainty, and applicable SGP interpretation. A structural indicator may directly support an existing structural or enabling-condition cell only under that cell's declared construct; it does not thereby establish experienced welfare.

Signed-impact contract. In I , positive means improvement and negative means harm relative to the declared baseline; zero means assessed no material change, not unknown. A nonnegative loss magnitude L , raw physical state z , susceptibility parameter χ , and signed impact I are different typed quantities. Converting among them requires an explicit, versioned mapping. Record units and sign conventions at the interface; do

not silently interpret positive shock magnitudes as positive welfare or clamp unsupported state values into apparent safety.

Observation and intervention records SHALL remain distinct and linked. Desired recovery, issued commands, resource allocations and successful message delivery do not establish an observed beneficial outcome. Corrections append to the evidence trail rather than overwrite a harmful observation. Existing qualification and authority rules apply to recovery actions as to any other consequence-bearing option.

2.2 Unions as the Unit of Analysis (Operational Definition)

Definition: A union is a bounded pattern of interdependence: a set of entities whose internal interactions are sufficiently strong, frequent, or consequential that their welfare should be evaluated together for the decision context.

In this specification, the operational objects scored and indexed by u are Union Scopes (Scopes), meaning unions used as stable scopes of aggregation and accountability.

Unions are analytical constructs for causal and welfare accounting; they are not metaphysical substances. The union concept operationalizes the systems-theoretic observation that complex systems exhibit hierarchical modularity (Simon, 1962).

2.2A Union Scopes and Instance Modeling (Normative)

Terminology clarification. Introduced in RippleLogic v12.0 and carried forward in v12.2, the canonical term Union Scope refers to the objects indexed by $u \in \{1, \dots, 7\}$ in the welfare matrix. For readability and reduced ambiguity, Scope is the default shorthand in technical and implementation contexts. The term Union remains an allowed legacy shorthand in narrative contexts and lineage documents. This clarification SHALL NOT change the mathematical meaning of u or any cascade semantics. Machine-readable surfaces SHOULD prefer UnionScope / UnionScopes as the canonical identifier family; Union remains a permitted human-readable alias in prose and UI text.

Purpose. Union Scopes are designed to be minimal but covering, a small stable set of aggregation scopes that prevents systematic stakeholder erasure while supporting auditability and comparability across decisions.

2.2A.1 Definitions (Normative)

Union Scope (Scope). A Union Scope is a stable scope of impact aggregation used to evaluate decision consequences across layered and intersecting analytical scopes. Union Scopes are the rows of the canonical 7×7 welfare matrix.

Stakeholder. A stakeholder is any entity that is plausibly materially affected by a candidate option a , including entities affected through externalities, indirect pathways, or future propagation within the declared decision horizon.

Instance. An instance is a stakeholder bound to the decision context as stakeholder-in-role-in-context, meaning a stakeholder with a declared role and exposure pathway relevant to the decision. Examples include: employee, patient, local resident, watershed ecology unit, cross-border supply-chain worker, digital agent affected by policy. Instances are the primary unit of stakeholder logging for auditability, Union Scopes are the primary unit of aggregation for comparison.

Structural anchor. Scopes $\times$ Dimensions form the welfare matrix. Instances feed the cells.

Local Stakeholder Set (LSS). The LSS is the context-specific set of stakeholder instances materially affected by a decision. It is the upstream representation object produced by stakeholder discovery and instance-to-scope mapping, and it is the justified input used to support any PLSS prominence declaration.

2.2A.2 Stakeholder Discovery Protocol (SDP) (Tier-linked; Normative)

Tier 3: REQUIRED. Tier 2: RECOMMENDED by default and REQUIRED when any SDP trigger holds. When no SDP trigger holds, a Tier 2 run MAY satisfy this requirement with a lightweight stakeholder note stating that no trigger held and why, or MAY voluntarily include a full SDP record. SDP is a short, explicit process for discovering omission-likely stakeholders so discovery is not optional by taste.

Minimum SDP steps (recorded in PCC): (i) define decision boundary and horizon, (ii) map the causal footprint and externality pathways, (iii) scan for rights exposure and catastrophe relevance, (iv) draft the Local Stakeholder Set (LSS) as the stakeholder instance set, (v) map instances to one or more Union Scopes with rationale, (vi) record unknowns and declared blind spots, including escalation triggers, (vii) provide a challenger opportunity for missing instances at Tier 3.

SDP triggers (Tier 2 REQUIRED): rights-covered cells plausibly affected, or TRC triggered/catastrophe relevance plausible, or externalities beyond decision owner plausible (cross-community/organization/polity, biosphere pathway), or irreversibility/lock-in plausible, or credible challenger claims omitted stakeholder class or scope, or information/procedure/enforcement/ecology pathways plausible.

2.2A.3 Instance-to-Scope Mapping Requirement (Tier-linked; Normative)

Tier 3 requirement. A Tier 3 run MUST explicitly log a set of stakeholder instances and map each instance to one or more Union Scopes, with a rationale sufficient for an independent reviewer to understand why the instance belongs in the mapped scope(s). This mapping SHALL be recorded in the PCC (Appendix H).

Tier 2 requirement under triggers. When any SDP trigger holds, Tier 2 runs MUST log and map stakeholder instances for the affected pathways and rights-relevant domains. Otherwise, Tier 2 runs SHOULD still log and map instances whenever the decision plausibly affects stakeholders beyond the decision owner.

Minimum mapping rule. Every logged instance MUST map to at least one Union Scope. If an affected stakeholder cannot be mapped, the PCC MUST record this as an explicit coverage limitation, not as neutrality, including: (i) why mapping was infeasible, (ii) what scope(s) might plausibly apply, and (iii) what risk this creates.

2.2A.4 Multi-scope Mapping and Redundancy Handling (Normative)

Multi-scope mapping is permitted. An instance MAY map to multiple Union Scopes where the decision plausibly affects that instance across multiple nested scales.

No silent duplication. An underlying effect MUST NOT acquire additional RLS mass merely because several scopes describe it. Declare one redundancy method: (i) EMERGENT_SCALE_JUSTIFICATION only for separately evidenced scale-specific DISTINCT_EMERGENT_OUTCOME tokens under Section 2.2A.4A, with token identities and challenger-visible evidence; the same effect is not repeated as a new emergent benefit; (ii) DEDUPLICATION, one PRIMARY_COMPUTATIONAL_HOME with VIEW_ONLY_NO_ADDITIONAL_VALUE tags, while separately evidenced downstream effects may be represented through the declared pathway model; or (iii) ALLOCATION, the ALLOCATED_CONSERVED_EFFECT treatment for one conserved cost, resource, exposure, transfer or burden, with non-duplicating coefficients under the following rule. The allocation exception distributes conserved accounting mass; it does not authorize copies of the same welfare outcome.

ConservedAllocationTokenID (allocation only). A ConservedAllocationTokenID is a redundant conserved causal unit used only to prevent redundant counting across scopes, for example the same dollars, the same one-time transfer, the same compliance cost. If ALLOCATION is used for a ConservedAllocationTokenID, the PCC MUST declare the token, the allocation method, the coefficients used, and the scope list. Allocation coefficients for a single ConservedAllocationTokenID MUST satisfy $\sum_u \alpha_u \leq 1$ for that token, summed over the scopes listed for that token.

No allocation laundering. ALLOCATION is permitted only for declared redundant conserved-unit effect-tokens. Distinct welfare harms, for example health harm versus ecosystem degradation, MUST be logged as distinct impact instances and SHALL NOT be allocated away.

Allocation guard. Allocation SHALL NOT be used to suppress or mask adverse impacts in any rights-covered or catastrophe-relevant evaluation. Redundancy handling cannot weaken NCRC, TRC, or CSV computations.

2.2A.5 Reduced-scope Mode (Normative)

Reduced-scope mode allows an evaluator to run with fewer than 7 scopes for mapping or reporting depth, provided omissions are explicitly disclosed. Reduced-scope mode SHALL NOT remove any non-maskable cells or required scope rows from NCRC, TRC, or CSV computation.

Reduced-scope mode is a reporting or mapping-depth choice only. Admissibility computation MUST still evaluate all rights coverage sets `C_r` and catastrophe cell set `C_cat` as declared.

If reduced-scope mode is used, the PCC MUST include a Scope Coverage Declaration (Appendix H) listing active scopes, omitted scopes, omission rationales, a blind-spot statement, and escalation triggers that would force scope expansion in the next comparable run.

Common misread (informative). Misread: “Unions limit who matters.” Correction: “Union Scopes prevent impacts from disappearing. Any stakeholder can be represented as an instance mapped into relevant scope(s), and omissions must be disclosed as coverage risk.”

2.2A.4A Effect Identity, Granularity, and Computational Home (Normative)

EffectTokenID. Every decision-material consequence SHALL have one stable EffectTokenID binding the affected bearer or bearer class, welfare or state construct, option, baseline, causal pathway, maximal contiguous decision-relevant time window, scenario treatment, and governing evidence object. Scope tags, labels, report views, or stakeholder interests do not by themselves create new welfare mass.

Token-individuation test. Two descriptions SHALL be treated as the same effect token when they refer to the same bearer/class, state change, baseline, causal pathway, contiguous time window, and evidence object. A distinct token is permitted only when at least one of those fields differs materially and the difference corresponds to a separately evidenced downstream or emergent consequence. Thorough description is not sufficient reason for finer tokenization.

Evidence-revision clarification. The governing evidence object is a stable provenance reference with a revision history, not a requirement to mint a new EffectTokenID whenever evidence is updated. Preserve the token when the represented consequence is unchanged; version its evidence and recompute affected conclusions. A materially different bearer, state, pathway or time window requires the individuation test, not automatic duplication.

Primary computational home. Each effect token SHALL have one primary Welfare Dimension and one primary aggregating cell for RLS. Secondary dimensions or Union Scopes MAY be recorded as `VIEW_ONLY_NO_ADDITIONAL_VALUE`. A separately evidenced downstream state MAY receive its own token and primary home.

ConservedAllocationTokenID. Where one conserved cost, resource, exposure, transfer, or burden is allocated across scopes, allocation coefficients SHALL sum to no more than 1.0. Any residual SHALL be explained. Allocation MAY improve visibility but MUST NOT suppress rights, catastrophe, severe-burden, or structural review.

RightsEffectTokenID. A decision-material rights harm SHALL also carry a RightsEffectTokenID and is governed by Section 7.4.5. Rights evaluation uses the unallocated token and cannot be attenuated by analytical scope allocation.

Granularity challenge. Tier 3 and high-stakes Tier 2 runs SHALL record whether a competent challenger proposed a materially different tokenization. Where two evidence-supported tokenizations remain and can change ranking or decisiveness, the PartitionInvarianceTest in Section 5.5A is mandatory.

Scope treatment	Meaning	RLS effect
PRIMARY_COMPUTATIONAL_HOME	One cell owns the token's aggregating contribution	Aggregates once
VIEW_ONLY_NO_ADDITIONAL_VALUE	Secondary visibility/accountability tag	No additional RLS mass
ALLOCATED_CONSERVED_EFFECT	One conserved token distributed with coefficients	Total allocation ≤ 1
DISTINCT_DOWNSTREAM_OUTCOME	Separately evidenced new state/consequence	New token may aggregate
DISTINCT_EMERGENT_OUTCOME	Scale-level state not reducible to the lower-level token	New token only with evidence and challenge

2.3 The Union Stack (Seven Operational Unions)

Union	Name	Definition	Characteristic Timescale
U ₁	Self	Individual-level affected stakeholder instances: each locus of experience and agency represented at the person-level scope. The decision-maker or operator is included only when materially affected and mapped like any other stakeholder; U ₁ is not a privileged decision-owner interest.	Seconds to Decades
U ₂	Household	Primary care, dependency, cohabitation, or sustained resource-and-support unit (operational definition).	Days to Decades
U ₃	Community	Recurring relational network with trust, identity, belonging or sustained interaction; may be geographically distributed (operational definition).	Months to Generations
U ₄	Organization	Formal collective pursuing a purpose; structured coordination and institutional behavior (March & Simon, 1958; North, 1990)	Years to Centuries
U ₅	Polity	Governance authority unit over a jurisdiction; legitimacy and institutional structure (Weber, 1978)	Decades to Centuries
U ₆	Humanity / Global Coordination	Humanity view: all affected human beings and future human populations. CMIU view: the governance-coordination network of intelligences, representatives, and institutions that have a declared participation channel or lawful role in shared governance, coordination, and systemic-risk management. The views may overlap but are not identical. CMIU is not a moral-status class: MPS or FPP does not create membership, and GPR does not create power. Direct participation uses a governed GPR record where relevant; consequential roles require a role-specific SPR record, lawful mandate, and all applicable RippleLogic gates.	Generations to Millennia
U ₇	Biosphere	Earth's integrated life-support systems, including climate stability and ecosystem integrity (Odum, 1971; Steffen et al., 2015; Rockström et al., 2009, 2023)	Centuries to Epochs

Canonical analytical scope sequence: $U_1 < U_2 < U_3 < U_4 < U_5 < U_6 < U_7$

The symbol $<$ denotes increasing analytical scope, not universal strict set inclusion. The seven rows are a coverage discipline for systematic review. Real applications MUST model material many-to-many stakeholder memberships, cross-border institutions, overlapping communities, and cross-scope causal dependencies through explicit stakeholder and dependency graphs.

U6 dual-view clarification. U6 remains one row coordinate so the 7×7 matrix stays stable. The **Humanity view** covers all affected human beings and future human populations. The **Collective Managing Intelligence**

Union (CMIU) view is the governance-coordination view of intelligences, representatives, and institutions that have a declared participation channel or lawful role in shared governance, coordination, and systemic-risk management. These views overlap but are not identical. CMIU is not a sentience tier, rights plateau, prestige class, or automatic membership granted by moral protection. FPP secures protection; GPR governs participation mode; SPR governs eligibility for a named consequential role; institutional membership and authority additionally require lawful mandate, domain competence, accountability, non-domination, auditability, and revocability.

U6 view-tag requirement (Normative). Every claim-bearing U6 input, impact, indicator, weight, scenario, or conclusion MUST declare U6_View_Tag as one of U6_HUMANITY, U6_COORDINATION, or U6_MIXED_WITH_DECOMPOSITION. A mixed U6 value is admissible only when the Humanity and coordination components are separately recorded and the aggregation rule is disclosed. U6 view tags prevent human welfare claims from being silently blended with institutional or managing-intelligence capacity claims.

Clarifying distinction (informative). Organization refers to bounded institutions and role systems such as firms, schools, NGOs, or agencies. Polity refers to public governance and enforcement domains such as lawmaking, courts, regulators, and due process systems.

Ecological trophic cascades illustrate how interventions can propagate unpredictably across levels and why “local wins” can generate system losses, reinforcing the need for explicit containment and ripple modeling (Estes et al., 2011).

2.3A Scope Interpretation and Coverage Continuity (Normative clarification)

The operational row identifiers remain U1-U7. Family/Household is a public expansion of the U2 Household label and covers primary care, dependency, cohabitation or sustained resource-and-support units; cohabitation or biological kinship is not mandatory. The run declares the actual unit and prevents duplicate accounting of the same care or transfer effect. U3 Community may be geographically distributed when recurrent interaction, trust, identity or belonging is evidenced. U4 Organization concerns purpose-directed role/rule systems; U5 Polity concerns public or jurisdictional governance, adjudication and enforcement. Formal rules, exit costs, network dominance or a protocol fork alone do not establish polity status or legitimate authority. Multi-role entities use typed, many-to-many membership and dependency records, not forced reassignment. U5 includes de jure and de facto public-governance domains where jurisdictional adjudication or enforcement is exercised; analytical classification is not a grant of legitimacy.

U6 remains Humanity / Global Coordination, with Collective Managing Intelligence Union (CMIU) as its explicit coordination view. CMI is permitted as a public abbreviation of that view, not a replacement for all affected or future human populations and not a separate eighth row. Preserve U6_HUMANITY, U6_COORDINATION and U6_MIXED_WITH_DECOMPOSITION, including mixed-view decomposition. Intelligence, sentience, protection, participation and authority remain separate. Nonhuman or digital entities can be represented at any materially applicable scope under the existing stakeholder and SGP rules; they are not confined to U6.

U7 Biosphere denotes Earth's integrated life-support systems, not the universe or a light cone. Solar, orbital and other external conditions may enter a bounded run through explicit causal dependencies and domain evidence without being relabelled as biospheric welfare. Material consequences outside the operational scope cannot be silently absorbed or omitted: disclose the coverage limit and narrow, escalate or invoke a governed extension. Cosmic and Universal/AIU remain the existing non-computational meta-union horizons by default. Row order is analytical, not proof of strict set containment.

2.4 Union Types vs Instances (Aggregation Rule)

The seven unions are types; real runs involve instances (many households, many organizations, many polities). Impacts MUST be aggregated within a union row using a declared aggregation method.

Default aggregation (normative default for non-rights scoring): Population-weighted mean across affected instances.

Rights exception: For rights-covered cells, apply worst-off subgroup checks within instances before aggregating (Section 7).

Distributional companion record (Normative supplement). When a non-rights welfare impact is materially concentrated, Tier 3 and high-stakes Tier 2 runs MUST report the worst materially affected subgroup or decile, the observed or estimated range, an inequality-direction indicator, and a burden-concentration ratio or declared equivalent alongside the population mean. These diagnostics do not automatically create a new gate, but they inform CSV, RLS sensitivity, stakeholder review, and whether the impact should be reclassified as rights-relevant.

2.4A Distributional Burden, Subgroup Closure, and Mandatory CSV Routing (Normative)

When the Distributional Companion Record is required, the record SHALL include the stakeholder/subgroup search method, worst materially affected subgroup, incidence, severity, duration, recurrence, reversibility, remedy, voice/appeal pathway, burden-concentration ratio, inequality direction, omitted-group uncertainty, challenger findings, and a SubgroupClosureStatus.

Allowed SubgroupClosureStatus values are SEARCH_DOCUMENTED, CHALLENGED_AND_REVIEWED, MATERIAL_OMISSION_OPEN, and INSUFFICIENT_SUBGROUP_GROUNDING. MATERIAL_OMISSION_OPEN or INSUFFICIENT_SUBGROUP_GROUNDING blocks a unique framework-selection claim when the omitted group could change RF/NCRC, TRC, CSV, RLS ranking, or decisiveness.

The Distributional Companion Record MUST be evaluated inside CSV against forced externalization, invisible or unmonitored harm, inability to consent/contest/recover, reversibility and remedy, burden accumulation, constitutive mitigation, and the existence of a materially cleaner or safer feasible alternative. A favorable population mean MUST NOT be the sole basis for clearing a severe concentrated burden.

Rights-adjacent D6 Meaning and local D7 Environment harms. In Tier 3 and high-stakes Tier 2 runs, a reviewed non-materiality determination or a governed rights-extension/routing record is required when a severe D6 or local D7 harm could implicate dignity, liberty, information, procedure, ecological integrity, bodily integrity, or basic needs. This rule does not automatically convert every D6/D7 burden into a rights violation.

Examples: Stakeholder Instances Mapped into Union Scopes (Non-Normative)

This table illustrates how diverse stakeholders are represented as instances mapped into Union Scopes. The scope set is minimal but covering, it is a coordinate system for accountability, not a closed list of who matters.

Stakeholder instance (example)	Typical role/context	Mapped scope(s) (example)	Notes (including multi-scope redundancy method when relevant)
Me (decision owner)	direct personal impact	U1 Self	Direct experience and agency locus.
Child or dependent in my care	care, safety, needs	U2 Household (and U1 for the child-as-person in detailed runs)	Household stability plus rights exposure often relevant.
Local neighborhood residents	local externalities	U3 Community	Trust, cohesion, local health and environment.
Employees or contractors	workplace policy impacts	U4 Organization	Pay, safety, dignity, due process.
Regulators, courts,	compliance and legitimacy	U5 Polity	Law, procedure, rights enforcement capacity.

Stakeholder instance (example)	Typical role/context	Mapped scope(s) (example)	Notes (including multi-scope redundancy method when relevant)
enforcement bodies			
Cross-border supply-chain workers	indirect externalities	U6 Humanity / Global Coordination	Captures externalities beyond a single polity.
Watershed or river system	ecological integrity	U7 Biosphere	Often ecology-rights and TRC relevant.
Future generations	long-horizon impacts	U6 Humanity / Global Coordination + U7 Biosphere	Multi-scope mapping. Default DEDUPLICATION with one primary home and view-only tags; separately evidenced scale-specific emergent outcomes require their own tokens. Horizon and indicators must be explicit.
Public information ecosystem	misinformation, censorship, legitimacy	U5 Polity + U6 Humanity / Global Coordination	Multi-scope mapping. If represented in both scopes, declare redundancy handling. Default is DEDUPLICATION with separately warranted propagation; EMERGENT_SCALE_JUSTIFICATION requires separately evidenced scale-specific tokens.
Digital agents affected by policy	AI deployment governance	U6 Humanity / Global Coordination (SGP-gated for protections where applicable)	Protections handled via the SGP interface. Authority remains separately gated.

2.5 Constructive vs Pathological Unions

Some unions grow by degrading their containing unions (for example, extractive industries or corruption networks). RippleLogic therefore does not assume “union benefit” is automatically good. The CSV Gate prevents local optimization that damages containing union coherence or viability beyond governed tolerance.

2.6 Meta-Unions (Non-Computational by Default)

U₈ Cosmic and U₉ Universal / All-Encompassing Infinite Union (AIU) may be used as philosophical boundary conditions or future extensions, but do not participate in standard Tier 1-3 scoring unless formally activated through a future governed extension protocol recorded in the canon, release manifest, and PCC requirements for that release line.

2.6A Operational UBR vs Metaphysical AIU (Normative Boundaries and Scope Discipline)

This section clarifies the strongest science-disciplined interpretation of Union-Based Reality (UBR), Union-Based Ethics (UBE), and the Universal/AIU horizon within the RippleLogic Canon v13.0 line. Its purpose is to preserve conceptual clarity, prevent overclaiming, and maintain a clean boundary between what is currently operational, what is normatively stipulated, and what remains philosophical, phenomenological, or future-facing.

Operational UBR. For the purposes of Tier 1 through Tier 3, UBR SHALL be interpreted operationally as a systems ontology of layered and intersecting causal interdependence. In governance-relevant domains, entities and decisions are embedded in coupled structures of consequence such that actions propagate across multiple scopes of impact and accountability. This is a descriptive claim about consequence topology, not a proof that all beings are identical in substance or ultimate ontology.

Scope discipline. Implementations and public claims MUST distinguish among: (i) the operational-descriptive layer, where consequence pathways are modeled across the layered and intersecting analytical scopes; (ii) the normative layer, where UBE enters through explicit commitment; and (iii) the meta-union horizon layer,

where Cosmic and Universal/AIU function as philosophical, humility, precautionary, or future-extension concepts. These layers MUST NOT be collapsed into one another.

Operational boundary. The standard scored union stack for Tier 1 through Tier 3 remains U1 Self, U2 Household, U3 Community, U4 Organization, U5 Polity, U6 Humanity / Global Coordination, and U7 Biosphere. U8 Cosmic and U9 Universal / All-Encompassing Infinite Union (AIU) are meta-unions and are non-computational by default in the current release line.

Permitted role of meta-unions. U8 and U9 MAY be used as philosophical boundary conditions, humility reminders, precautionary amplifiers, or future-extension placeholders. They MUST NOT, absent a governed extension protocol, be presented as standard Tier 1 through Tier 3 scoring objects, admissibility overrides, or evidence-bypass mechanisms.

Interpretation of “we are one.” Within RippleLogic Canon v13.0, statements of oneness SHALL be interpreted in three distinct registers: (i) an operational register of strong interdependence and multi-scope accountability; (ii) a structural-unity research register, concerning whether convergent mathematical or physical principles underlie the nested interdependence observed across scales, and whether such principles constrain viable governance architectures; and (iii) an absolute-metaphysical register that remains outside current computational proof status. Only the first register is part of standard Tier 1 through Tier 3 governance computation.

Phenomenology and horizon claims. Reports of unity, non-separation, or all-encompassing union arising from contemplative, mystical, near-death, psychedelic, or out-of-body experiences MAY be treated as phenomenological data, horizon signals, or motivation for future inquiry. They MUST NOT, by themselves, be treated as sufficient evidence for standard parameterization of U9 or for compliance claims under the canon.

Non-overclaiming rule for AIU language. The framework MAY be described as general-purpose and multi-scale, and as aiming toward broadly universalizable governance architecture for embedded intelligences in shared causal reality. It MUST NOT be described as having scientifically proven All-Encompassing Infinite Union, fully resolved all metaphysical disagreements, or already solved universal alignment in the strongest ontological sense.

Canonical summary. RippleLogic adopts UBR as an operational ontology of layered and intersecting interdependence within measurable shared reality, adopts UBE through an explicit normative entry, and retains Cosmic and Universal/AIU language as non-computational horizon layers unless and until a governed extension protocol makes broader parameterization legitimate.

SECTION 3: NORMATIVE FOUNDATIONS: UNION-BASED ETHICS (UBE)

3.0A Ethical Reality

Ethical reality is the subset of reality where structures, constraints, and consequences affect welfare-bearing beings and the conditions of flourishing. Ethical reality includes welfare, rights, sentience, agency, dignity, non-domination, catastrophic-risk exposure, ecological dependence, future generations, and flourishing. Access to ethical reality spans measured fact, uncertainty-bound inference, and explicit normative commitment. These layers MUST NOT be collapsed: measured indicators may support claims, SGP may govern uncertain moral-patient evidence, and rights floors operate as declared non-compensatory commitments rather than as empirical measurements alone. Ethical reality is therefore not separate from reality; it is reality where consequence matters to beings and systems capable of being harmed, protected, dominated, liberated, degraded, or enabled to flourish.

3.1 Minimal Normative Axiom (MNA)

RippleLogic adopts one minimal normative entry axiom, then operationalizes it through governed commitments, including rights floors, tail-risk constraints, containment and structural-viability requirements, constitutional weight floors, and auditability rules. These operational commitments are explicit governance priors, not hidden mathematical derivations. The minimal entry axiom is:

MNA: Sentient welfare and flourishing matter. Unnecessary suffering and domination should be reduced. The enabling conditions for viable, non-dominated, and continuing flourishing should be preserved.

This axiom is:

Content-minimal: No single “good life” doctrine is prescribed.

Scope-bounded: Applies where sentience or welfare-relevant moral patienthood is in scope.

Architecturally generative: The framework defines a compositional architecture, rights floor, tail-risk constraint, containment, welfare optimization, and structural safeguards, such that diverse domain-specific instantiations can be generated by declared inputs and then executed through the same cascade.

Normatively explicit: RippleLogic does not claim to derive this axiom from mathematics alone or from descriptive facts alone.

3.1A Core Definitions for the Normative Layer (Normative)

Interpretation rule: RippleLogic protects welfare-bearing subjects directly and protects enabling conditions structurally where their degradation predictably undermines welfare, rights protection, non-domination, or long-run system viability.

3.1B Bridge Premises (Normative support; non-deductive)

The MNA is accompanied by two bridge premises that reduce arbitrariness while preserving honesty about the entry of normativity.

Bridge Premise 1: Embedded Modelling. Any agent seeking reliable action in shared causal reality must model other stakeholders, including their interests, constraints, vulnerabilities, dependence relations, rights exposure, and, where behavior, harm, cooperation, or rights protection depend on them, welfare-relevant states and enabling conditions, to a non-trivial degree.

Bridge Premise 2: Prudential Stability. In plural, uncertain, co-embedded systems, architectures that systematically degrade stakeholders' welfare conditions, rights floors, or shared enabling conditions increase instability, conflict, tail risk, and long-run governance failure.

These premises do not generate obligation by deduction. They explain why welfare-tracking, rights-tracking, and enabling-condition tracking are structurally relevant to intelligence acting in shared reality.

3.1C Normative Commitment (Explicit final step)

RippleLogic adopts the commitment that credibly evidenced welfare-relevant states, rights-relevant states, and the enabling conditions on which they depend constitute reasons for action, not merely inputs for prediction.

This final step, treating such states and conditions as reasons rather than mere data, remains a normative commitment, not a deductive conclusion. RippleLogic holds that this is among the narrowest defensible gaps between description and obligation currently achievable in a governance-grade framework without either claiming false deduction or retreating to brute preference, and treats closing it further as an active research question.

3.1D Evidence Discipline for Normative Claims (Normative)

Within RippleLogic, a putative welfare-relevant state, rights-relevant state, or enabling condition SHALL enter governance computation only through declared evidence, governed interpretation, or an explicit precautionary treatment under uncertainty appropriate to the implementation tier.

If such a state or condition materially affects admissibility, precautionary handling, selection, or tie-break outcome, it MUST be recorded in PCC.NormativeInputRegister with sufficient provenance for independent review.

Unsupported narrative assertion, rhetorical force, linguistic fluency, or undeclared evaluator intuition SHALL NOT by itself constitute a sufficient reason-for-action input within RippleLogic.

3.2 Conditional Is-Ought Bridge

RippleLogic avoids deriving values from facts by using a conditional bridge:

IF

(UBR) actions propagate through layered and intersecting analytical scopes,

AND

(MNA) sentient welfare and flourishing matter, unnecessary suffering and domination should be reduced, and the enabling conditions for viable, non-dominated, and continuing flourishing should be preserved,

AND

(BP1) agents operating in shared causal reality must model other stakeholders' interests, constraints, vulnerabilities, dependence relations, rights exposure, and, where behavior, harm, cooperation, or rights protection depend on them, welfare-relevant states and enabling conditions to act reliably,

AND

(BP2) in plural, uncertain, co-embedded systems, systematic degradation of stakeholders' welfare conditions, rights floors, or enabling conditions increases instability, conflict, tail risk, and long-run governance failure,

THEN

agents and institutions ought to evaluate choices by cross-union impacts and preservation of enabling conditions, using non-compensable rights floors, explicit tail-risk bounding, containment safeguards, and auditable welfare ranking within the feasible set.

Normativity enters explicitly through the MNA and the final normative commitment. UBR specifies where consequences flow. BP1 and BP2 explain why welfare-tracking, rights-tracking, and enabling-condition tracking are structurally relevant to intelligence acting in shared reality. RippleLogic does not claim that the final normative step is fully deduced from description alone.

3.2A Derived Governance Rule (Normative)

In governance contexts, RippleLogic operationalizes the above commitment through the canonical lexicographic cascade:

Define the qualification predicates without collapsing claim authority into ethical rejection:

$RG_qualified(a) := 1$ iff $RG_status(a)$ in $\{RG_SUPPORTED, RG_NARROWED\}$. $RG_NARROWED$ means the option may be evaluated only inside the narrowed claim boundary. $RG_REFUSED$ means the present claim is not supportable at the requested strength and downstream ordinary evaluation short-circuits; it is not a claim that the option is intrinsically unethical in every possible reformulation.

$TRC_qualified(a) := 1$ iff $TRC_status(a)$ in $\{TRC_PASS, TRC_NOT_TRIGGERED\}$. $TRC_NOT_TRIGGERED$ is available only after a recorded catastrophe-relevance assessment finds no trigger under Section 8.1.1. It is

not a silent skip, does not mean CVaR was computed, and must reopen if new catastrophe relevance appears.

$A_sel := \{a \text{ in } O : RG_qualified(a) \text{ AND } NCRC(a) \text{ AND } TRC_qualified(a) \text{ AND } CSV_status(a) \text{ in } \{CSV_PASS, CSV_PASS_WITH_CONTROLS, CSV_NOT_MATERIAL\}\}$. This is the Phase 1 qualification set. If A_sel is empty, escalate, redesign, or invoke governed fallback per the canon. Otherwise Phase 2 ranks over A_sel using RLS and the declared tie-break chain.

1. Protect rights floors first through NCRC.
2. Bound catastrophic downside second through TRC.
3. Block pathological containing-scope degradation through CSV as the third option-filtering gate after the RG claim-authority precondition.
4. Rank remaining feasible options fourth through welfare scoring under declared weights, uncertainty rules, and auditability requirements.
5. Resolve non-decisive cases structurally through UCI/HOI tie-breaks and escalation rules.

Interpretation rule: RippleLogic does not define “ought” as unconstrained welfare maximization. It defines “ought,” in governance-grade decision settings, as selection from a rights-safe, tail-bounded, containment-safe feasible set under auditable evidence and uncertainty discipline.

3.2B Open Research Frontier (Informative)

An active research question remains whether sufficiently accurate multi-agent world-modelling creates deeper constitutive pressure toward non-domination and welfare protection beyond prudential stability alone. RippleLogic does not assume this as a proven result. It treats it as a frontier question concerning whether accurate modelling of shared reality may eventually narrow the remaining descriptive-to-normative gap further. Specifically, this question is tractable through empirical study of whether systems that model co-embedded agents more accurately also exhibit more stable cooperative and non-destructive behaviour, and whether omission of welfare-relevant, rights-relevant, or enabling-condition variables from world models degrades predictive performance in domains where cooperation, conflict, or harm depend on them.

3.3 Seven Welfare Dimensions (Canonical)

Dimension	Name	Description
D ₁	Material	Resources and infrastructure for survival and functioning
D ₂	Health	Physical and mental functioning; morbidity and mortality risk
D ₃	Social	Belonging, trust, relational integrity, cooperation
D ₄	Knowledge	Epistemic access and learning conditions
D ₅	Agency	Autonomy and effective choice; freedom from coercion
D ₆	Meaning	Coherence, purpose, valued life projects (measured cautiously)
D ₇	Environment	Ecological and built context integrity sustaining life

This seven-dimensional choice is a convergence architecture rather than a claim of metaphysical completeness. It aligns with the capability approach (Sen, 1999; Nussbaum, 2011), self-determination theory (Ryan & Deci, 2000), and fundamental human needs theory (Max-Neef, 1991), while also capturing ecological integrity as a life-support substrate consistent with Earth-system boundary research (Rockström et al., 2009, 2023).

7×7 field interpretation rule (Normative for claim language). The canonical 7 Union Scopes × 7 Welfare Dimensions matrix is a multidimensional accountability-and-flourishing field. Depending on the bearer and cell, a recorded condition may represent direct welfare, collective functioning, an enabling condition, structural integrity, ecological condition, epistemic observation or correction capacity, or a symbolic/meaning condition. The stable interface name Welfare Dimension is retained for compatibility, but a cell MUST NOT be described as phenomenally experienced welfare of a non-sentient system unless admissible evidence supports that claim. WDBIP records the applicable effect-token and cell-type interpretation and prevents determinants, observation capacities, or system conditions from being mislabeled as directly experienced welfare.

3.4 Non-Fungibility at Rights Level

Dimensions are treated as non-fungible at the rights layer: gains in one dimension MUST NOT compensate rights-floor violations in another. This is enforced structurally by NCRC (Section 7), not by weight tuning.

3.5 Unioning (Enacted Practice)

Unioning = redesigning decisions until they are rights-safe (NCRC), tail-safe (TRC), containment-safe, and net-positive across unions under declared uncertainty policy, treating apparent “tradeoffs” as design failures before acceptance.

Unioning worked example (Non-Normative)

Scenario. An organization considers cutting safety training to reduce costs. Initial estimate shows a small material benefit in U4 Organization D1, but a credible safety degradation in U4 D2 and a dignity/procedure risk in U4 D6/D5. Under NCRC, any rights-covered violation depth $v_r(a) > 0$ fails admissibility regardless of aggregate gains.

Step 1, fail. Option A fails NCRC due to increased probability of serious injury and due-process erosion. It is rejected even if it would raise RLS.

Step 2, redesign. The team introduces Option B, maintain safety training, add targeted efficiency changes elsewhere, and add a transparent incident-reporting and remediation pathway. This shifts the harmful impacts above the rights floor while preserving most of the intended cost reduction through non-harmful channels.

Step 3, pass and select. Option B passes NCRC, passes TRC and CSV, then competes on RLS among admissible options. The PCC records the redesign rationale and the rights-floor reasoning, so the improvement is legible and replayable.

3.6 NCAR Learning Loop (Corrigibility Requirement)

RippleLogic is embedded in Notice, Choose, Act, Reflect. The system is corrigible: it records assumptions and outcomes, updates kernels and parameters via governed procedures, and preserves accountability through versioned PCCs.

SECTION 4: SYSTEM OVERVIEW: THE RIPPLELOGIC ARCHITECTURE

Public method. RippleLogic SHOULD be taught and implemented at the public level as:

Reality Grounding -> Rights Floor -> TRC -> CSV -> RLS

Human mnemonic:

Ground -> Protect -> Bound -> Contain/Verify -> Score

Formal method. In formal notation and audit surfaces, this corresponds to: RG/RSG -> RF/NCRC -> TRC -> CSV -> RLS, where RG is Reality Grounding and RSG is the reality-surface grounding implementation record.

Clarification. Reality Grounding is a claim-authority precondition, not an ethical option-rejecting gate.

RG_REFUSED still excludes an option from the claim-bounded qualified set until grounding is repaired or the claim or action boundary is narrowed. RF/NCRC, TRC, and CSV gate ethical or structural admissibility and selectability. RLS ranks the selectable set. The five-level method therefore consists of one claim-authority precondition (RG), three hard feasibility/selectability gates (RF/NCRC, TRC, CSV), and one residual ranking layer (RLS).

Two-phase decision model. The five-level method may be taught and implemented as two phases without changing the cascade. **Phase 1 qualifies options:** RG establishes claim authority; RF/NCRC protects rights; TRC bounds ruin; CSV verifies containment and structural viability. **Phase 2 ranks survivors:** RLS orders only the options that remain selectable after Phase 1.

Qualification-before-ranking rule. Do not score fantasy, rights violations, ruin paths, or structurally non-viable options as ordinary alternatives. RLS is applied only to the selectable set produced by RG/RSG, RF/NCRC, TRC, and CSV.

TRC/CSV non-collapse rule. TRC is the ruin veto: it asks whether catastrophic, irreversible, lock-in, systemic, or ruin-path downside is acceptably bounded. CSV is the viability and control test: it asks whether an option can structurally stand, execute, remain contained, preserve dependencies, maintain monitoring, and remain governable under real constraints. CSV may discover a TRC-relevant scenario, but it does not absorb TRC. If CSV identifies a new ruin path, TRC MUST be reopened before RLS.

Level 1, Reality Grounding (RG/RSG). Declare the reality surface, evidence trace, material unknowns, transition boundary, consequence pathways, category boundaries, and claim boundary. If grounding is insufficient for a material claim, narrow, escalate, mark exploratory, or refuse the stronger claim.

Level 2, Rights Floor (RF/NCRC). Apply the Non-Compensatory Rights Constraint. Rights-floor harms are non-compensatory and cannot be rescued by welfare score.

Level 3, TRC / Tail-Risk Bound. Apply the Tail-Risk Constraint. Catastrophic or unbounded downside cannot be averaged into ordinary welfare gains.

Level 4, CSV / Containment and Structural Viability. Apply the CSV gate. CSV asks whether the option's local gain can stand without unacceptably damaging, hollowing, overloading, exploiting, or depending on the larger systems that make the gain possible. UCI/HOI and related structural diagnostics are evaluated inside CSV when material. Structural Viability is explicit inside CSV: resource closure, dependency closure, endpoint/reversibility coherence, operational capacity non-exceedance, and internal non-contradiction.

Level 5, RLS / RippleLogic Score. Rank only the options that remain selectable. If RLS is tied, close, or non-decisive, record any governed residual UCI/HOI preference under Section 11.5 while retaining REFUSE_DETERMINISTIC_SELECTION unless the Section 10.4/10.4B unique-selection conditions are met.

Mantra. Ground -> Protect -> Bound -> Contain/Verify -> Score.

Layer	Question answered	Output
RG / Reality Grounding	What reality surface, evidence trace, unknowns, transition boundary, consequence pathways, category boundaries, and claim boundary support this run?	Grounded, narrowed, escalated, exploratory, or refused.
NCRC	Does the option violate a non-compensatory rights floor?	Rights-admissible or rejected.
TRC	Is catastrophic tail risk acceptably bounded?	Admissible or rejected.
CSV	Can the option deliver its claimed good without uncontained harm, hollowing, lock-in, structural non-viability, or degrading containing systems beyond tolerance?	Selectable, pass with controls, redesign required, emergency provisional, excluded, or escalated.
RLS	Among selectable options, which has the strongest residual welfare pattern?	Ranked selectable options.

4.1 Formal Definition

RippleLogic is a rights-first, tail-risk-bounded, union-based decision optimization system with explicit ripple propagation, auditable scoring, and corrigible learning. It is designed to prevent three common governance failures: rights being traded away by aggregate welfare gains, catastrophic downside being averaged away by expected-value logic, and local optimization degrading larger containing systems.

4.2 Core Objects and Implementable Data Structures

Every RippleLogic run works from declared objects. These are not decorative labels; they are the minimum data surfaces that make the decision auditable, reproducible, and challengeable.

Object	Description
Option set O	Finite set of candidate actions.
Welfare impact matrix per option	$I_{\text{prop}}(u,d,a)$ in $[-1,+1]$, after direct impact construction, propagation, and saturation. Cell interpretation is governed by the canonical 49-cell reference grid and Appendix AD.
Rights impacts	$I_{\text{rights}}(u,d,a)$, using worst-off protected subgroup semantics for rights-covered cells.
Scenario set S with probabilities p_s	Scenario surface for TRC, including mandatory tail categories and probability floors where required.
Catastrophe cells C_{cat} and weights ω	Catastrophe-relevant welfare cells and tail-risk weights used in TRC.
Kernel K, sparse 49x49	Propagation structure. NONE or QUICK modes are permitted in Tier 1-3 subject to the canon.
Weights w_u and v_d	Union-scope and dimension weights, governed by HDW where required and constitutional floors where applicable.
Structural metrics UCI and HOI	Containment and tie-break diagnostics, subject to maturity boundaries.
ReferenceStructureRecord	The declared reference object, evidence status, causal pathway, falsification condition, authority boundary, and refusal trigger for the claim being made.

4.3 Reference Structure Discipline

RippleLogic decisions SHALL NOT be interpreted as governance claims merely because a scoring procedure, audit record, or policy wrapper exists. A governance claim is supportable only to the extent that the decision record declares the reference structure against which the claim is evaluated.

For RippleLogic purposes, a reference structure is the declared set of constraints, evidence sources, causal pathways, thresholds, maturity labels, and refusal conditions used to evaluate a decision claim. Reference structures may include physical constraints, biological needs, rights floors, catastrophic-risk corridors, containment conditions, stakeholder impact maps, institutional authority boundaries, evidence-status labels, and auditable records.

RSP and RefStructRecord. The Reference Structure Protocol (RSP) is the discipline. The ReferenceStructureRecord (RefStructRecord) is the run-level artifact that records the discipline in a specific case. RSP tells the system what must be declared; RefStructRecord records what was actually declared.

Reference structure discipline does not create a sixth gate and does not alter the canonical method. Reality Grounding, NCRC, TRC, CSV, RLS, and UCI/HOI diagnostic handling retain their existing roles. The Reference Structure Protocol strengthens the evidence basis used to evaluate those roles.

A RippleLogic run MUST NOT treat circular, self-validating, missing, unfalsifiable, materially stale, or unsupported reference claims as sufficient for high-stakes admissibility, selection, deployment-readiness, or conformance claims. Where the reference structure is insufficient for the requested claim, the correct outcome

is Refusal, escalation, narrower claim scope, additional evidence gathering, or bounded provisional action under an applicable triage rule.

Plain-language rule. Governance without a declared reference structure becomes assumption management. RippleLogic requires declared references, visible evidence, falsifiable claims, and refusal when the reference structure cannot support the claim.

4.4 Formal Gate Sequence and Residual Tie-Break (Normative)

RippleLogic evaluates candidate options through a formal gate sequence with a claim-authority precondition. Reality Grounding first establishes whether the requested claim can be made at the declared strength. RF/NCRC, TRC, and CSV then determine admissibility and selectability. RLS ranks only the selectable set. Residual UCI/HOI is a special-case tie-break or monitoring rule, not a regular public cascade level. A later function MUST NOT rescue, compensate for, or reinterpret failure at an earlier function.

Gate-numbering clarification (Normative). Formal gate numbering counts only the three hard admissibility/selectability gates: RF/NCRC, TRC, and CSV. RG is the upstream claim-authority precondition and is not numbered as a hard gate; RLS is the downstream ranking stage.

Formal gate 1, NCRC: rights feasibility. NCRC determines whether an option is rights-admissible. For rights-covered cells, RippleLogic applies the required subgroup semantics and checks the worst-off protected subgroup against the rights floor. Any option with positive violation depth on any covered right fails NCRC and is inadmissible, regardless of expected welfare gains, downstream rankings, or structural advantages.

Formal gate 2, TRC: catastrophe feasibility. TRC determines whether a rights-admissible option is admissible in the tail-risk sense. It evaluates catastrophic-loss exposure over the governed scenario set using the catastrophe cell set and declared CVaR corridor. Any option whose catastrophic tail exceeds the corridor fails TRC and is inadmissible, even if expected welfare performance is strong.

Formal gate 3, CSV: containment and structural viability feasibility. CSV determines whether an admissible option is selectable with respect to containing-system integrity and execution viability. It asks whether local or sub-scope welfare gains degrade containing unions, hollow out capacities, create lock-in, externalize burdens, exceed operational capacity, or fail structural viability. An option that fails CSV is not selectable, even if it passes NCRC and TRC and would otherwise rank highly on residual welfare.

Ranking stage, RLS: residual welfare ranking. RLS is applied only to options that have already passed RF/NCRC, TRC, and CSV. It ranks surviving options by weighted welfare aggregation over the active cell set, subject to non-maskable-cell rules and weight-governance requirements. RLS MUST NOT override or soften any earlier gate.

Special-case UCI/HOI tie-break rule: structural resolution after RLS non-decisiveness. UCI and HOI are first evaluated inside CSV when material to structural integrity, hollowing risk, or containment. Only after options are already selectable may residual UCI/HOI help document a tied, close, uncertainty-overlapped, or non-decisive RLS result. UCI and HOI do not act as public cascade stages and do not reopen prior gate decisions.

Non-overlap rule. For interpretation and auditability, the formal gate/ranking functions SHALL be read as follows: RG/RSG = claim authority; NCRC = rights feasibility; TRC = catastrophe feasibility; CSV = containment and structural viability feasibility; RLS = residual welfare ranking within the selectable set; UCI/HOI tie-break = special-case non-decisive structural resolution or monitoring within the selectable set, not a regular public cascade level. No implementation, commentary, or downstream artifact may collapse these functions into a single blended score.

Level	Name	Function	Failure consequence
1	RG / Reality Grounding	Claim grounding and reality-surface discipline; claim-authority precondition, not an ethical option-rejecting gate.	Narrow, escalate, mark exploratory, or refuse stronger claim; RG_REFUSED is outside the qualified set.

Level	Name	Function	Failure consequence
2	Rights Floor (RF/NCRC)	Rights feasibility.	Rejected except Emergency Mode.
3	TRC / Tail-Risk Bound	Catastrophic tail-risk feasibility.	Rejected or escalated.
4	CSV / Containment and Structural Viability	Containing-system integrity, structural viability, hollowing/dependency/lock-in screening, and controls/redesign logic.	Pass, pass with controls, redesign required, emergency provisional, excluded, escalated, or refused.
5	RLS / RippleLogic Score	Residual welfare ranking among selectable options.	Selection if decisive; otherwise disclose refusal and any governed tie-break preference separately from authority selection (Section 11.5).

4.4A Structural Viability within CSV (Normative clarification)

Structural Viability is explicit inside CSV, not a sixth public method level. It asks whether an option that otherwise passes Reality Grounding, NCRC, TRC, and containment integrity can structurally execute within the declared reality surface without undefined resource, dependency, reversibility, operational-capacity, or internal-contradiction failure.

Formal shorthand: RG/RSG -> RF/NCRC -> TRC -> CSV -> RLS. Where material, the CSV result includes Structural Viability status before any option enters RLS. An option may be ethically preferable as a redesign target while still failing current structural viability for execution. In that case, the option may be documented as SVC_DESIGN_TARGET or CSV_REDESIGN_REQUIRED, but it MUST NOT be claimed as currently selectable, executable, deployment-ready, or deterministically selected.

Non-import boundary. Structural Viability uses only MathGov-native records: resource declarations, dependency declarations, reversibility and endpoint status, operational limits, and contradiction checks. It MUST NOT import capacity/binding ontology, regime coordinates, hidden kernels, or undisclosed reference generators.

4.5 Cascade Gate Inputs and Outputs (Informative)

This table summarizes what each gate consumes, what it produces, and what it may eliminate or flag. It is a fast correctness check for implementers.

Gate / layer	Consumes	Produces	Eliminates / flags
Reality Grounding layer	Reality surface, evidence trace, material unknowns, transition boundary, consequence pathways, claim boundary, and reviewer/evidence status where material.	RealityGroundingStatus; claim boundary; refusal/escalation/narrowing disposition.	Blocks, narrows, escalates, marks exploratory, or refuses claims that exceed the grounded surface; does not compare options as an admissibility gate.
NCRC rights floor	Rights-specific floor-channel impacts $I_{RF}(u,d,a g,r)$ and $I_{rights}(u,d,a,r)$ for	RF_PASS, RF_FAIL, RF_ESCALATE or NCRC_UNKNOWN, with channel and subgroup evidence	A pass requires every applicable rights channel to pass with no unresolved material field. Missing or contested material evidence cannot be collapsed into a Boolean zero.

Gate / layer	Consumes	Produces	Eliminates / flags
	C_r; subgroup operator gamma where governed fallback is required; rights thresholds; categorical and severe-hazard records.		
TRC tail-risk / CVaR	Scenario set S with categories and floors; catastrophe loss proxy L(a,s) over C_cat; alpha; tau_TRC.	TRC_PASS, TRC_NOT_TRIGGERED, TRC_FAIL or TRC_ESCALATE, with screen or loss evidence	Apply the Section 8.6 predicate, including the active severity overlay. TRC_NOT_TRIGGERED requires its own supported trigger assessment.
Containment system safety	Containment triggers; structural posture thresholds; Delta UCI / HOI / Adaptive Reserve or equivalent diagnostics where material; reversibility and repair evidence.	One of the seven CSV statuses in Section 9.2, with controls, uncertainty and remediation records	Eliminates or escalates options failing containment; flags CONTAINMENT_STRUCTURAL_DIAGNOSTIC_UNAVAILABLE or CONTAINMENT_UCI_HOI_UNAVAILABLE when required evidence is missing.
RLS welfare selection	Active welfare cells, including rights-covered and catastrophe cells subject to non-maskable rules; weights w_u and v_d; applicability mask m(u,d); uncertainty proxies.	RLS(a); Gap(a,b); welfare ranking over remaining options.	Does not override NCRC/TRC/CSV; flags invalid masking of rights or catastrophe cells from RLS aggregation.
Tie-break / monitoring diagnostics	Residual UCI/HOI or equivalent structural indicators not already decisive for Containment; monitoring and hollowing documentation.	Tie-break signal, monitoring warning, or future-run steering.	Not a gate; may be used only when RLS is tied, close, uncertainty-overlapped, or non-decisive, or for monitoring. Cannot rescue earlier failure.

Binary predicates used elsewhere are derived acceptance tests over the full status record, not substitutes for that record. Unknown, escalation, not-triggered and not-evaluated states retain their explicit meanings.

4.6 End-to-End Algorithm (Tier 1-3 Executable) (Normative)

Algorithm: RippleLogic_Run

Inputs, minimum: decision scope; baseline; option set O; union and dimension sets; weights; rights canon; TRC canon and scenario set wherever the applicable tier and trigger require them; containment parameters; propagation mode; kernel if used.

Inaction-asymmetry rule (Normative). When no action, delay, or continuation of the current course is a live feasible choice, it MUST appear as an explicit candidate option and receive its own Reality Grounding, rights, tail-risk, CSV, and effect-matrix treatment. It MUST NOT be represented only as a null baseline exempt from evaluation. The PCC SHALL disclose material continuation harms and the reason any apparently available no-action or delay path was excluded.

Outputs: selected option a* or escalation, plus PCC.

Step 1: Reality Grounding and Notice. Record the decision question, scope boundary, time horizon, affected unions, and option set. Declare tier, configuration, reality surface, evidence trace, material unknowns, transition boundary, consequence pathways, claim boundary, and reference-structure requirements. If grounding is insufficient for a material claim, narrow, collect evidence, escalate, mark exploratory, or refuse the stronger claim before proceeding.

Step 2: Impact construction. For each option, construct the declared Base and Welfare impact instances and uncertainty treatment. Under NONE, $I_{prop} := I_{dir}$ without a second saturation. Under QUICK, apply the declared propagation and post-propagation saturation per scenario before probability averaging where scenario conditioning applies. Gate-specific non-attenuation and missing-data rules remain controlling.

Step 3: NCRC gate. Evaluate the floor, categorical and severe-hazard channels, worst-off subgroup treatment and unallocated rights-effect rules in Section 7.4.4. A pass requires every applicable violation to be zero and no unresolved material field. If no option passes, apply the Section 7.4 rights-conflict precondition and seek redesign or governed escalation. Rights Emergency Mode may be considered only after all Section 7.5 prerequisites, including independent challenge, are satisfied.

Step 4: TRC gate. For each rights-admissible option, perform the required trigger screen or scenario assessment. A computed pass requires the full Section 8.6 predicate, including CVaR and every active severity-class condition. If no option remains, return TRC_ALL_FAIL and use redesign, delay/no-action comparison, escalation or refusal. Tail Emergency Mode may be considered only when every Section 8.8 prerequisite is demonstrated; it never emits an ordinary TRC_PASS.

Step 5: CSV gate. For each NCRC/TRC-admissible option, apply the containment predicate using required containing-scope, structural posture, reversibility, repair, and evidence-availability inputs. Where UCI, HOI, Adaptive Reserve, lock-in, hollowing, trust erosion, ecological degradation, governance-capacity loss, or related structural diagnostics are material to cross-scale structural integrity, evaluate those diagnostics inside CSV before RLS. If required structural diagnostic evidence is unavailable, downgrade, collect evidence, rerun, escalate, refuse the stronger claim, or record a permitted waiver where the applicable tier permits waiver.

Step 6: RLS ranking. Compute RLS(a) only for selectable options. Compute uncertainty and apply the discrimination rule where required. If decisive, select the top option.

Step 7: Special-case tie-break / non-decisive handling. If RLS is tied, close, uncertainty-overlapped, or otherwise non-decisive, compare only options already in the selectable set. Apply residual UCI/HOI only as tie-break, monitoring, or hollowing-risk documentation. Do not reopen NCRC, TRC, or CSV. Do not use residual UCI/HOI to rescue failed Containment or compensate for missing containment evidence. A tie-break preference does not establish RLS decisiveness: retain REFUSE_DETERMINISTIC_SELECTION while Section 10.4/10.4B remains unmet, and apply Section 11.5 serialization before any separate AuthoritySelectionRecord.

Step 8: Emit PCC. Record inputs, intermediate computations, gate results, sensitivity outputs where required, ReferenceStructureRecord where required, refusal record if applicable, and public rationale.

4.7 Canonical 49-Cell Reference Table (v10.6+; Informative)

Each cell in the 7×7 welfare matrix represents a distinct accountability domain. The table below provides a compact one-line meaning for each cell to support comprehension and consistent interpretation across implementations. The table is informative; the governing equations and rules remain in the relevant normative sections and appendices.

Interpretive note. This table gives the compact canonical grid. Appendix AD provides the expanded 49-cell Welfare Dictionary, including plain meanings, positive and negative movement, indicator examples, possible gate relevance, examples, and reviewer checks. Appendix AD does not change the cascade, equations, rights thresholds, TRC mechanics, Containment semantics, RLS formula, HDW, SGP, PCC, RSP/RefStructRecord, or claim boundaries. It governs scoring interpretation and reviewer literacy.

Union / Dimension	D1 Material	D2 Health	D3 Social	D4 Knowledge	D5 Agency	D6 Meaning	D7 Environment
U1 Self	personal resources	personal health and safety	personal relationships and cohesion	personal knowledge and learning	personal agency and autonomy	personal meaning and purpose	personal environmental conditions
U2 Household	household resources	household health and safety	household relationships and cohesion	household knowledge and learning	household agency and autonomy	household meaning and purpose	household environmental conditions
U3 Community	community resources	community health and safety	community relationships and cohesion	community knowledge and learning	community agency and autonomy	community meaning and purpose	community environmental conditions
U4 Organization	organizational resources	organizational health and safety	organizational relationships and cohesion	organizational knowledge and learning	organizational agency and autonomy	organizational meaning and purpose	organizational environmental conditions
U5 Polity	polity/public resources	polity/public health and safety	polity/public relationships and cohesion	polity/public knowledge and learning	polity/public agency and autonomy	polity/public meaning and purpose	polity/public environmental conditions
U6 Humanity / Global Coordination	humanity-wide resources	humanity-wide health and safety	humanity-wide relationships and cohesion	humanity-wide knowledge and learning	humanity-wide agency and autonomy	humanity-wide meaning and purpose	humanity-wide environmental conditions
U7 Biosphere	biospheric resources	biospheric health and safety	biospheric relationships and cohesion	biospheric knowledge and learning	Resilience Capacity	Life Continuity	biospheric environmental conditions

U7 scope-specific operationalization rule (Normative clarification). U7/D3 Biotic Relations, U7/D4 Ecological Knowledge, U7/D5 Resilience Capacity and U7/D6 Life Continuity use declared ecological or enabling-condition analogues, not personal social experience, conscious knowledge, autonomy or experienced meaning. Use the Appendix AD definitions and a typed construct/bearer record. Human ecological knowledge and institutional monitoring have a primary human or institutional D4 home, with U7 view-only tagging; they cannot receive a second U7/D4 contribution merely because their subject is ecology. A separate U7 information-bearing condition needs separate evidence and a distinct token. U7/D5 records shock absorption and recovery; U7/D6 records continuity of life and lineages. U7/D7 may represent an integrated life-support condition, but must not add the same underlying state again after its component D2/D3/D5/D6 tokens have been scored. An integrated TRC view may reuse linked evidence without adding a second RLS contribution or masking catastrophe coverage. Retain the 7 x 7 interface, declare each construct exception, and test material alternative cell treatments or weights. Ecological analogue scores are not personal-welfare construct-validation data.

Release packaging rule (v10.6+). Every public release PDF MUST include this 49-cell table and the stakeholder-to-scope examples table below, either in Section 4 or in a dedicated reference-table appendix that is directly linked from Section 4.

Scenario	Primary scope logging	Secondary / cross-scope impacts	Redundancy handling note
Parent transfers \$100 to adult child.	Recipient: U2 Household (child	Sender household reduced resources, U2	Distinct bearer effects: record the recipient's gain and the sender's loss separately, each with one primary home

Scenario	Primary scope logging	Secondary / cross-scope impacts	Redundancy handling note
	household) x D1 Material.	x D1; governance/admin impacts logged separately if applicable.	and a shared transfer provenance link. DEDUPLICATION removes repeated views of the same bearer's consequence, not the other party's loss. Monetary conservation does not imply equal welfare effects.
City spends \$1M on public hospital upgrades.	Community: U3 Community x D2 Health.	Households benefit; polity accountability impacts may exist.	Allocate only a conserved investment cost or resource quantity, with disclosed coefficients. Record separately evidenced health, access and accountability outcomes as distinct tokens; do not split a non-conserved welfare benefit merely across reporting scopes.
National carbon tax: distinguish revenue and compliance costs from evidenced ecosystem-state and human-welfare consequences.	Record ecosystem-state effects at their appropriate U7 homes and distinct population effects at their justified human-scope homes, with U6 view tags where used.	Polity administrative burden and compliance costs logged as distinct impacts.	ALLOCATION applies only to a conserved revenue, cost, resource or burden token with coefficients summing to at most one. Do not allocate an undifferentiated policy-benefit score; evidence each distinct outcome.
Factory emissions affect local residents and ecosystem.	Local health harm to residents: U3 Community x D2 Health.	Ecosystem degradation: U7 Biosphere x D7 Environment.	Distinct downstream or bearer-specific tokens require distinct states and evidence. A scale label alone is not EMERGENT_SCALE_JUSTIFICATION.
River communities 50km downstream affected by discharge.	Downstream residents: U3 Community, separate instance from factory-neighbor community.	Potential rights exposure flagged through stakeholder discovery.	SDP externality detection: omission-likely stakeholder class logged separately.

4.8 Implementation Tiers (Normative Summary)

Requirement	Tier 1 Heuristic	Tier 2 Core	Tier 3 Auditable
NCRC	Heuristic	REQUIRED	REQUIRED with subgroup semantics
TRC	Required qualitative trigger screen	REQUIRED when catastrophe relevance is plausible	REQUIRED
Scenario set size	N/A	5 minimum	20 minimum
Kernel propagation	NONE	NONE default; QUICK if KQS policy satisfied	QUICK only if KQS policy satisfied plus sensitivity
CSV qualification	REQUIRED before any ordinary ranking/selection claim; qualitative screen or CSV_NOT_MATERIAL rationale allowed	REQUIRED proportional/assumption-bounded CSV result before RLS	REQUIRED full and binding
PCC	Optional	REQUIRED basic	REQUIRED full
Sensitivity analysis	Optional	Recommended	REQUIRED

Tier 4 boundary. Tier 4 is a design target only in the current v13.0 release and MUST NOT be claimed until ProofPack is public and independently replayable.

4.9 Refusal Under Underdetermination (Normative)

Purpose. RippleLogic distinguishes gate failure from underdetermination. Gate failure occurs when a declared option uniquely violates NCRC, TRC, or CSV requirements. Refusal occurs when the framework cannot uniquely determine whether an option passes, fails, or can be selected from the declared inputs and projection rules.

Refusal is not weakness, discretion, or rhetorical caution. It is a formal anti-overclaiming state. A run that enters Refusal MUST NOT claim that the relevant option is admissible, selectable, selected, aligned, safe, or RippleLogic-conformant for the unresolved claim.

4.9.1 Refusal triggers (Normative)

R-1 Gate-critical unknown. A rights-covered cell, catastrophe cell, Containment indicator, or other gate-critical input remains UNKNOWN_IMPACT after the applicable evidence, phantom-instance, or conservative-default rule has been applied.

R-2 Non-unique projection. The projection from stakeholder instances, subgroup evidence, scenario data, or scope mapping to a gate verdict admits two or more materially different admissibility interpretations under the same declared evidence.

R-3 Material indistinguishability. Two or more options are observationally indistinguishable on available welfare evidence but materially different in possible rights exposure, tail-risk exposure, domination risk, or containment risk.

R-4 Layer ambiguity. The run relies on a category collapse, including treating compliance as correctness, traceability as admissibility, psychological plausibility as evidence, mathematical elegance as validation, moral status as governance authority, or SGP welfare scalars as admissibility modifiers.

R-5 Projection modification after observation. A ProjectionPreRegistration object, threshold, rights coverage set, catastrophe cell set, scenario set, weight rule, SGP MPS/FPP binding, or Containment trigger is materially modified after relevant decision evidence has been consulted, unless the run is voided and restarted.

R-6 Claim-boundary mismatch. The evidence available may support a lower-tier, local, provisional, or exploratory claim but is insufficient for the stronger claim being attempted.

4.9.2 Refusal effects (Normative)

1. stop any claim of admissibility, selectability, selection, conformance, or alignment for the unresolved option or run;
2. emit a PCC.RefusalRecord;
3. identify the unresolved gate, scope, dimension, scenario, stakeholder class, or projection rule;
4. state the minimum evidence or procedural condition required for re-run;
5. route the decision to escalation, redesign, bounded provisional action, or no-action baseline review according to tier and stakes.

4.9.3 Emergency and provisional action (Normative)

Refusal blocks strong conformance claims. It does not automatically imply inaction in emergency contexts. Where delay itself creates material harm, a provisional action MAY proceed only if all of the following hold:

1. the action is rights-preserving under the best available conservative screen;
2. catastrophic tail-risk is not plausibly increased beyond the declared emergency corridor;
3. the action is reversible, staged, or bounded where feasible;
4. escalation and independent review are triggered;

5. the PCC clearly marks the action as `PROVISIONAL_UNDER_REFUSAL` rather than admissible, selectable, or selected under normal conformance.

4.9.4 Distinction from gate failure (Normative)

Gate failure is committed: NCRC violation, TRC corridor breach, or CSV failure produces inadmissibility or non-selectability under unique projection. Refusal is uncommitted: the framework cannot uniquely determine admissibility from the declared inputs. Implementations **MUST NOT** silently convert Refusal into gate failure, gate success, or RLS tie-break, and **MUST NOT** use Refusal to bypass any earlier gate.

4.9.5 PCC fields for Refusal (Normative)

The `PCC.RefusalRecord` **SHALL** include at minimum: trigger code R-1 through R-6; evidence basis with provenance; options affected; scopes and dimensions affected; rights, catastrophe, or containment cells affected; no-action baseline status where relevant; escalation route invoked; provisional-action status, if any; and conditions under which a re-run could resolve the Refusal.

4.9.6 Relation to ripple.md IND and NA (Informative)

In ripple.md wrapped deployments, IND maps to Refusal where the test applies but evidence, scope lock, target specification, constraint set, verification regime, or projection conditions are insufficient. NA means the test is genuinely not applicable under the declared scope. Implementations **MUST NOT** use NA where IND is the correct state.

4.10 Consequence-tempo and responsibility-continuity interface (Normative routing)

Execution assurance must distinguish decision admissibility from the speed and ownership of implementation. For actions involving rapid propagation, material irreversibility, consequential automation, or material lock-in, the worst credible control critical path must be shorter than the time to an unacceptable adverse escalation threshold or irreversible stop-loss boundary. Where no runtime interruption window is technically possible, stronger ex ante assurance, bounded scope and rate, fail-safe design, and explicit residual-risk authorization are required. Delay must also be tested as a possible source of rights, ruin, or structural harm.

Complexity may distribute work but must not dissolve answerability. Claim-bearing high-consequence runs **SHALL** preserve typed responsibility for evidence, analysis, decision authority, operation, execution, intervention, appeal, communication, remedy, and residual responsibility. Legal liability is not inferred unless separately established by law.

4.10A Material obligation integrity (Normative)

Where selectability, execution, review, appeal, intervention, monitoring, rollback, communication, repair, or requalification depends on a material obligation, the PCC **SHALL** include or reference a `MaterialObligationRecord`. The record **SHALL** identify: the obligation; accountable authority; operational carrier or authenticated execution path; authority and capacity basis; activation trigger; scope; required action; response deadline or window; evidence of discharge; review or expiry; challenge route; delegation acceptance; backup or successor; nonperformance trigger; missed-duty escalation; amendment, waiver, suspension, and retirement authority; change-control rule; residual responsibility; and control-effectiveness status.

A declaration, committee, register, policy, review date, or named owner is not by itself a binding control. The carrier must have the authority, competence, information access, time, resources, and protection necessary to act. A handoff is incomplete until the receiving role accepts it. Failure, absence, conflict, incapacity, or silence by the carrier **SHALL** activate the declared escalation, substitution, pause, or requalification path. A material change to the control's trigger, threshold, scope, carrier, authority, evidence surface, expiry, escalation path, waiver rule, or retirement rule is a configuration change and **SHALL** reopen affected qualification.

Control performance and control effectiveness are distinct. Evidence that a procedure was followed does not establish that the intended right, risk, containment, or repair objective was achieved. Where effectiveness is decision-material, the run SHALL record the outcome evidence or remain assumption-bound, controlled, escalated, or redesigned.

4.10B Hidden human compensation load (Normative inside CSV when material)

Before automation, scaling, role removal, or claims of structural resilience, CSV SHALL test whether functioning materially depends on undocumented human compensation, including repeated manual override, exception reconciliation, shadow records, individual memory, relationship-based coordination, after-hours recovery, weak-signal interpretation, informal escalation, or recurrent prevention of visible failure. Such compensation is evidence of a structural dependency; it is not by itself evidence that the system is resilient.

When material, the PCC SHALL include or reference a HumanCompensationLoadRecord identifying the carrier role classes, hidden work performed, formal-workflow gap, frequency or burden indicators, single-person dependency, recovery and capacity effects, consequence if the compensation stops, evidence sources, privacy and non-retaliation protections, preservation or redesign action, and review trigger. Evidence SHOULD be triangulated from workflow traces, overrides, exceptions, rework, handoffs, near misses, workload and recovery indicators, and protected testimony rather than depending solely on self-report or waiting for carrier collapse.

These are cross-cutting assurance requirements routed through TRC and CSV and recorded through ripple.md v5.8 and Agent System v13.0. They are not a sixth gate and do not make a completed record proof of safety.

SECTION 5: WELFARE IMPACT CONSTRUCTION (CALCULABLE)

5.0 Purpose and Design Requirements (Normative)

This section specifies how RippleLogic converts real-world predictions and evidence into a computable welfare impact matrix for each option. A competent analyst MUST be able to compute all required impacts from (i) a declared baseline, (ii) a declared option, and (iii) a finite list of impact instances per cell, using only the equations and rules in this paper and appendices.

Cell-meaning rule (Normative clarification). A cell score is a reviewable baseline-relative claim, not a preference expression. It SHALL be interpretable as an assertion about one Union Scope and one welfare dimension, with evidence and uncertainty disclosed according to tier requirements. Appendix AD provides the canonical interpretive dictionary for cell meanings, positive/negative movement, possible indicators, gate-relevance cues, examples, and reviewer checks.

5.0.1 Stakeholder Coverage Index (SCI) (Normative)

SCI is a tier-governed disclosure ladder describing the stakeholder coverage posture of a run. SCI does not change scoring formulas. It is an auditability and integrity signal that supports comparability and anti-theater enforcement.

SCI levels (minimum definitions): SCI-0: no explicit stakeholder instance logging. SCI-1: instances logged at least for direct stakeholders and primary externality pathways. SCI-2: instances logged with explicit mapping to Union Scopes, including at least one challenger pass for omissions. SCI-3: evidence-backed instance mapping with explicit blind-spot declaration, escalation triggers, and redundancy handling disclosures where multi-scope mapping is used. SCI-4: high-coverage posture with structured discovery, third-party review or validated stakeholder registry linkage, and replay-ready traceability.

Tier minima. Tier 2 runs MUST achieve at least SCI-1. Tier 3 runs MUST achieve at least SCI-2.

Anti-theater ratchet. If SCI equals the tier minimum, the PCC MUST include a concrete Next-Run Upgrade action that would raise SCI by at least one level for the next comparable run. If SCI is below the tier minimum, the run MUST either (a) downgrade the declared tier to one whose SCI minimum is satisfied, or (b) rerun with

enhanced stakeholder discovery to meet the SCI minimum for the claimed tier. In either case, the PCC MUST include a remediation plan specifying how SCI will be raised to the target level before the next comparable run. A remediation plan alone does not make the current run compliant at the originally claimed tier.

5.0A End-to-End Cell Computation Recipe (Informative)

This recipe is the minimal end-to-end spine for computing any welfare-cell impact in a Tier 1-3 run. It is informative, but the equations and constraints it references are normative in their home sections. For traceability, this recipe is listed in reverse order (8 → 1), reflecting the final artifact emission back to upstream inputs.

8) Emit PCC: record parameters, ReachBasis, gate results, RLS summary, and the public 5SPR. If reduced-scope mode is used, include the required attestations and blind-spot statement.

- 7) Run the formal method: RSG / Reality Grounding → RF / Rights Floor (formal NCRC) → TRC / Tail-Risk Bound → CSV / Containment and Structural Viability → RLS / RippleLogic Score. Treat UCI/HOI or equivalent structural diagnostics as CSV evidence when material, and as residual tie-break/monitoring evidence only when RLS is tied, close, uncertainty-overlapped, or non-decisive.

6) Apply kernel propagation only if declared: if KernelPropagation = QUICK, compute $I_{prop_welfare}(u,d,a)$ per Section 6; else set $I_{prop_welfare} = I_{dir}$. (Redundancy handling such as EMERGENT/DEDUP/ALLOCATION is a separate instance/aggregation mechanism and MUST NOT be treated as a propagation mode.)

5) Compute the direct contribution and accumulate to the cell: $\tilde{I}_{dir}(u,d,a;h) = \sum_k (\mu_k \cdot r_k \cdot \tau(t_k) \cdot \ell_k \cdot c_k \cdot e_k \cdot h_k)$ over k in that cell. Apply the canonical saturation (for example tanh) to obtain $I_{dir}(u,d,a;h) \in [-1,+1]$. Here h_k is a declared welfare-inclusion hypothesis coefficient, not an MPS score or moral-worth scalar.

Stream note (Normative). For admissibility layers (NCRC, TRC, CSV), compute the Base stream with $h_k=1$ for all impact instances. In the Welfare stream used for RLS, MPS evidence determines the required protection posture and hypothesis set, but the MPS band or interval MUST NOT be inserted directly as a cardinal multiplier. Human persons and FPP entities use $h_k=1$. For a non-FPP stakeholder whose welfare-bearing status is decision-material, the run MUST evaluate the governed inclusion hypotheses required by Section 12 and Appendix G, including $h_k=1$ and, where exclusion remains evidentially admissible, $h_k=0$. If the decision changes across the hypothesis set, the run returns MPS_HYPOTHESIS_SENSITIVE and may not conceal the sensitivity with a convenient intermediate value.

4) For each impact instance k , declare its parameters inside canonical bounds: magnitude $\mu_k \in [-1,+1]$, reach $r_k \in [0,1]$ with ReachBasis, time horizon $t_k \in (0, \infty)$ years (used via $\tau(t_k)$), likelihood $\ell_k \in [0,1]$, confidence $c_k \in [0.1, 1]$, bounded adjustment multiplier $e_k \geq 0$ (default 1.00; non-default use governed by Section 5.2 notes and PCC disclosure), and welfare-inclusion hypothesis $h_k \in \{0,1\}$ unless a separately governed and validated welfare-attention profile authorizes additional bounded values. h_k is never inferred mechanically from an MPS interval. Every human person and FPP entity uses $h_k=1$ (Appendix G).

Dual-baseline implementation note (Normative clarification). In Appendix P, every gate-relevant cell retained in RLS uses the FLOOR_REFERENCE value in both the gate and welfare streams. A different STATUS_QUO welfare baseline for any such cell is permitted only through a governed BaselineDualExceptionBlock that states the rationale, both values, non-compensation effect, sensitivity result, authority, and review trigger. In the absence of that block, floor-reference remains controlling. For ordinary non-gate welfare cells, STATUS_QUO may remain the declared welfare baseline.

Note on evidence (Normative clarification): Evidence quality is recorded in PCC provenance fields and informs the choice of c_k . Evidence quality is not a separate multiplicative scalar in $\tilde{I}_{dir}$ unless explicitly introduced by a future canonical revision.

3) Construct Impact Instances k for each active welfare cell (u,d) : each k must reference (a) a stakeholder instance, (b) a cell (u,d) , and (c) a direct or indirect pathway description.

2) Build the representation layer: SDP → Stakeholders → Stakeholder Instances (InstanceMap with union-scope coverage).

1) Declare the run: option set A, unions U_run, welfare dimensions D, horizon window, and kernel propagation policy (NONE / QUICK). If effect-token redundancy handling is used, declare the redundancy mode (EMERGENT / DEDUP / ALLOCATION) and α budgets separately.

5.0B WDBIP welfare-evidence interface (Normative when WDBIP conformance is claimed)

The Welfare Dimension Boundary and Interaction Protocol v1.9 is the Canon-subordinate companion for primary dimension assignment, token-level welfare pathways, dependence and cross-scope duplication disclosure, time-window and subgroup integrity, measurement uncertainty, and welfare-dimension validation planning.

WDBIP does not alter the seven dimensions, Union Scopes, gate sequence, RLS equation, constitutional weights, or lawful authority. It cannot rescue failed RF/NCRC, TRC, or CSV. When WDBIP conformance is claimed, the exact WDBIP record, schema version, and content hash SHALL be attached to or immutably referenced from the PCC.

The Canon owns effect-token definitions, Union Scope mapping, redundancy methods, allocation constraints, RLS weights, and gate routing. WDBIP consumes those objects by version and pointer and MUST NOT redefine them.

5.0C Distinct Effects, Accounting Homes and Representation Limits (Normative)

One effect, one computational home is an accounting discipline, not a theorem that real systems are independent. An effect keeps a stable identity across repeated observations, replicated messages, document views and evidence updates. A new source document, identifier, sensor reading or scope tag alone does not create a new consequence. Preserve both event/observation identity and effect identity; record revisions and causal parent links. Distinct downstream or emergent outcomes require distinct state/bearer/horizon meaning and their own warranted contribution.

The seven dimensions are designed for conceptual distinguishability and low representational redundancy while preserving real causal coupling. Do not claim strict orthogonality, statistical independence, causal independence, proven maximality, or an irreducible/exhaustive seven-part ontology without an applicable definition and evidence. Seven is a governed, revisable coverage architecture. Measurement correlation alone neither proves duplication nor justifies merging two causally related constructs.

Computational-home and conserved-allocation rules do not guarantee invariance under arbitrary repartitioning, changed weights or nonlinear saturation. Preserve the unallocated rights effect; secondary views add no RLS mass. Alternative admissible tokenizations or allocations that could change the score or decisiveness require the existing partition and robustness review. A one-time cost and the outcomes it enables must not both proxy the same final benefit without an explicit endpoint-accounting justification.

5.1 Baseline-Zero Rule (Normative)

Semantic anchor (MUST): For all unions u , dimensions d , and options a :

For all unions u , dimensions d , and options a , $I(u,d,a) = 0$ denotes zero net represented change from the declared baseline at the cell level.

Interpretation: “0 means no net represented change from baseline in the cell” is globally enforced. Any method that produces absolute levels MUST immediately convert them into baseline-relative deltas before entering RippleLogic impact scoring.

Interpretive boundary (MUST): A zero cell value does not establish that every underlying indicator, stakeholder subgroup, time window, or causal pathway is unchanged. Material opposing contributions MUST remain separately disclosed in the impact-instance, subgroup, distributional, and interaction records and MUST NOT be netted away for RF/NCRC, TRC, CSV, subgroup, or peak-harm analysis.

Stream binding (Normative): Compute a Base stream and a Welfare stream. The Base stream fixes $h_k := 1.0$ and is used by TRC and CSV. RF/NCRC begins from the Base stream but MUST apply the rights-specific non-attenuation, categorical-prohibition, and severe-hazard rules in Section 7.4; adverse rights severity MUST NOT be reduced merely because duration is short or analyst confidence is low. The Welfare stream uses the governed welfare-inclusion hypothesis set defined by the SGP binding interface and is used only for RLS. MPS bands and intervals are evidence objects and protection triggers, not cardinal welfare weights. If ranking or selection changes across the hypothesis set, the PCC MUST report MPS_HYPOTHESIS_SENSITIVE and the run must seek evidence, apply a precautionary policy, narrow the claim, preserve multiple options, or refuse a definitive ranking.

5.1A Baseline Contract for Admissibility Layers (Normative)

Baseline scope (Normative). RippleLogic welfare impacts are baseline-relative by design (Baseline-Zero Rule). For admissibility layers (NCRC and TRC) and any containment checks that depend on gate-relevant cells, the baseline for any cell that is in a rights coverage set C_r or in the catastrophe set C_{cat} MUST represent a rights-safe / safe-corridor reference condition (the rights-reference baseline; legacy machine alias FLOOR_REFERENCE), not merely the status quo.

Floor-baseline requirement (Normative). For each active rights or catastrophe cell, the PCC MUST declare the floor baseline reference (indicator definition, target level, jurisdictional or governed source when available, and any conversion assumptions). An option that maintains a rights-violating status quo MUST therefore score below 0 on the gate-relevant baseline-delta scale, because the baseline for gating is the floor reference.

PCC disclosure (Normative). The PCC MUST record `BaselineType_Gates = FLOOR_REFERENCE` for all cells in $C_r \cup C_{cat}$, and MUST record `BaselineType_Welfare` for other cells. Allowed values: `STATUS_QUO`, `FLOOR_REFERENCE`, `OTHER_DECLARED`. If `OTHER_DECLARED` is used, the basis and rationale MUST be specified and MUST NOT weaken NCRC/TRC protection.

For catastrophe cells, the safe-corridor baseline remains the controlling TRC reference. Section 8.4 fixes the scenario-conditioned residual-loss interpretation and separates it from the diagnostic same-scenario incremental comparison.

5.2 Impact Instances (Normative)

RippleLogic represents impacts in each active welfare cell (u,d) for option a as a finite set of impact instances $K(u,d,a)$. Each instance captures one distinct causal pathway or measurable effect.

InstanceMap vs Impact Instances Bridge (v9.0; Normative): InstanceMap (Section 2.2A) specifies who is represented (stakeholders-in-context). Impact instances $k \in K(u,d,a)$ specify how impacts occur (pathways) within a `Scope×Dimension` cell. These are distinct objects and MUST NOT be conflated. Impact instances inherit scope mapping from the stakeholder instances they affect. Both are required for full auditability at Tier 3.

Informative: `Scopes × Dimensions` form the welfare matrix. Instances feed the cells.

5.2.1 ReachBasis canon (Tier 2+ required) (Normative)

Definition (Normative). Reach $r_k \in [0,1]$ is a scale term representing the fraction of the relevant reference class affected by impact instance k for the specified union scope u and welfare dimension d . r_k is not a moral weight and MUST NOT substitute for rights floors, TRC, or other gates.

ReachBasisType (Normative). For each r_k , the PCC MUST record one `ReachBasisType` $\in \{\text{POPULATION_FRACTION, ASSET_OR_FLOW_FRACTION, AREA_OR_VOLUME_FRACTION,}$

OTHER_DECLARED}. If OTHER_DECLARED is used, the PCC MUST define the basis and the denominator unambiguously.

Denominator rule (Normative). The denominator MUST match the union scope u and the decision boundary (Section 2.2A), and MUST be disclosed with units (for example: adult residents in district X; monthly platform active users; hectares of watershed; annual procurement spend).

Estimator and bounds (Normative). The PCC MUST record the estimator method and source/provenance for the denominator and numerator, plus conservative bounds when uncertainty is material. If bounds are used, the PCC MUST record r_{k_low} and r_{k_high} and use the conservative bound in any gate-relevant analysis when reach uncertainty could affect admissibility or containment.

Sign-aware reach-bound rule (Normative). When reach uncertainty is material, the PCC MUST use the conservative endpoint relative to the claim being supported: a beneficial effect used to justify selection, safety, control adequacy, or mitigation uses the lowest supported reach; an adverse effect uses the highest supported reach; and a gate-critical adverse effect uses GateAdverseBound. If sign is mixed or unresolved, evaluate both supported endpoints. If admissibility, containment, ranking, or selection changes across the interval, record sensitivity and narrow, escalate, redesign, delay, or refuse the stronger claim. A denominator choice MUST NOT systematically reduce adverse reach or inflate beneficial reach.

Rights-specific reach semantics (Normative). For rights-covered impact instances used in NCRC evaluation, reach SHALL be defined relative to the materially exposed protected subgroup or another subgroup-exposed reference class that prevents attenuation by unaffected population fraction. Reach in rights-covered cells MUST NOT dilute severity by reference to the full union population when only a protected subgroup is exposed. The PCC MUST record the rights-specific denominator when it differs from the general welfare denominator.

Table 5.1: Impact Instance Parameters (Canonical)

Parameter	Symbol	Range	Meaning
Magnitude	μ_k	$[-1, +1]$	Signed severity of welfare change (baseline-relative). Positive = improvement; negative = harm.
Reach	r_k	$[0, 1]$	Fraction of the relevant stakeholder population materially affected.
Time horizon	t_k	$(0, \infty)$ years	Duration over which the effect remains materially relevant.
Likelihood	ℓ_k	$[0, 1]$	Conditional probability the instance occurs (conditional on scenario model, if used).
Confidence	c_k	$[0.1, 1]$	Analyst confidence in the instance specification and mapping (floor prevents “zeroing out”).
Equity/resilience multiplier	e_k	Tier 2: 1.00 only unless named governed instrument; Tier 3: $[0.75, 1.25]$ default 1.00	Governed multiplier for declared equity/resilience/context corrections only; MUST NOT function as a hidden weight or sign-changing control.
Welfare-inclusion hypothesis	h_k	Ordinarily 0 or 1; $h_k=1$ for every human person and FPP entity	Declared welfare-accounting hypothesis used for robust RLS sensitivity. It is not an MPS score, probability of consciousness, quantity of experience, or measure of intrinsic worth. Additional bounded values require a separately governed and validated welfare-attention profile.

Normative notes:

- c_k has a floor of 0.1 to prevent omission-by-zeroing. In the residual Welfare stream, c_k is a governed evidence-support shrinkage coefficient toward the declared no-material-change prior; it is not the complete uncertainty model and it is not a claim that the unshrunk effect is absent. The PCC MUST separately report uncertainty under Section 10.3 and MUST test whether ranking changes under a governed unshrunk or adverse-bound sensitivity when low-confidence adverse impacts are material. In RF/NCRC, TRC, and material CSV use, the Gate-critical confidence guard below controls: low confidence cannot operate as a protective discount.

Gate-critical confidence guard (Normative). For an adverse gate-critical instance k in a rights coverage set C_r , catastrophe cell set C_{cat} , or material CSV indicator, define $\text{GateAdverseBound}(k)$ deterministically before the run. If an admissible interval or scenario bound is available, $\text{GateAdverseBound}(k)$ is the most adverse supported endpoint for harm claims and the least favourable supported endpoint for beneficial control, capacity, or mitigation claims. If only a point estimate and c_k are available, the adverse gate calculation removes confidence shrinkage by setting $c_{gate,k}=1$, while a beneficial claim used to establish gate safety retains $c_{gate,k}=c_k$. Likelihood and scenario probability remain separately represented and MUST NOT be double-counted. If no reproducible conservative bound can be justified and the gate outcome could change, the result is UNKNOWN, ESCALATE, NARROW, or REFUSE rather than PASS. The PCC MUST record the bound method, evidence source, endpoint, and disposition. Confidence MUST NOT launder severe but uncertain gate-critical harm into pass status.

- e_k is tightly governed. Tier 2 MUST set $e_k = 1.00$ for all impact instances unless a named governed instrument from PCC.EKRegistry is invoked explicitly. Tier 3 MAY use non-default e_k only from a pre-declared registry of allowed instruments, with each instrument carrying a declared purpose, source, allowed range, and replayable reason code. Unless a stricter charter rule exists, allowed non-default values are bounded to $0.75 \leq e_k \leq 1.25$. e_k MUST NOT reverse sign, erase or dilute a rights-floor breach, function as a hidden weight, or be used asymmetrically across compared options without explicit justification and sensitivity disclosure. Any non-default e_k requires PCC justification, declared reason code, source or instrument note, stated bounds, default-counterfactual disclosure (what happens at $e_k = 1.00$), and tier-appropriate sensitivity if the non-default value could materially affect admissibility or selection.

Audit trigger (Normative). If a non-default e_k is used without declared rationale, source or instrument note, stated bounds, or the required sensitivity record for the claimed tier, implementations MUST set audit flag `EK_PARAMETER_UNJUSTIFIED`. `EK_PARAMETER_UNJUSTIFIED` is INVALID for Tier 2-3 claimability until corrected or downgraded.

Registry omission rule (Normative clarification). PCC.EKRegistry MAY be omitted when all impact instances in the run use $e_k = 1.00$ and no Tier-2 named-instrument exception is invoked.

Registry note (Normative clarification). Implementations SHOULD restrict non-default e_k uses to a small declared registry (for example `EQ_CORR`, `RESILIENCE_CORR`, `CONTEXT_CORR`, or `OTHER_DECLARED`) and MUST apply the declared rule symmetrically across compared options unless the asymmetry itself is the subject of evaluation.

Confidence-versus-uncertainty distinction (Normative clarification). The instance factor c_k is a declared governance shrinkage operator used to construct a conservative point estimate; it is not itself the uncertainty interval or probability that the effect is true. Cell-level $\sigma(u,d,a)$ represents a separate uncertainty surface. When the combined use of shrinkage and uncertainty could change admissibility, ranking, or decisiveness, the PCC MUST report an unshrunk effect sensitivity or an equivalent interval/distributional treatment and explain why the two operators are not double-counting the same evidence weakness.

- SGP v8.8 MPS records may affect protection posture, evidence requirements, and the welfare-inclusion hypothesis set used for robust RLS sensitivity. They MUST NOT be inserted directly as cardinal welfare multipliers. Every human person and FPP entity uses $h_k=1$. For other SGP-aware cases, include $h_k=1$ whenever direct welfare is a realistic or credible possibility; include $h_k=0$ as an exclusion hypothesis only where supported by admissible evidence. MPS-NE requires both inclusion and exclusion hypotheses when decision-material. SGP outputs MUST NOT weaken protection or alter NCRC, TRC, or CSV (Section 12; Appendix G).

5.2B Likelihood Semantics and Probability Ownership (Normative)

Every non-unit likelihood SHALL declare one LikelihoodSemantics value: `CONDITIONAL_EVENT_PROBABILITY`, `EXPECTED_FREQUENCY_OR_EXPOSURE`, `ALREADY_EXPECTED_EFFECT`, `SCENARIO_CONDITIONAL`, or `OTHER_DECLARED`.

`CONDITIONAL_EVENT_PROBABILITY` means the full consequence occurs with the declared probability and otherwise does not occur or is replaced by another scenario-defined state. It SHALL be evaluated scenario-first under Section 5.8 and Appendix B.

EXPECTED_FREQUENCY_OR_EXPOSURE means the multiplier is an expected incidence, duty cycle, exposure fraction, or repeated-rate term rather than an all-or-nothing event. Expected-input construction MAY be used when justified.

ALREADY_EXPECTED_EFFECT means the upstream magnitude already contains the relevant probability or frequency. The RLS likelihood multiplier SHALL then be 1.0 for that probability surface.

SCENARIO_CONDITIONAL means scenario probability p_s owns the occurrence weighting. The same probability MUST NOT also enter ℓ_k . Probability may be reused in TRC because TRC evaluates a separate catastrophe-loss object, but the provenance link and non-double-counting disposition SHALL be explicit.

OTHER_DECLARED requires a mathematical definition, unit/interpretation, aggregation order, source, reviewer, and double-counting check. It is not a license to choose the operator producing the preferred result.

If two likelihood interpretations are evidence-supported and selection-material, both SHALL be computed. A leader or decisive/non-decisive reversal sets PROBABILITY_ORDER_SENSITIVE and prevents an unconditional unique-selection claim.

LikelihoodSemantics	Correct mathematical home	Prohibited use
CONDITIONAL_EVENT_PROBABILITY	Scenario outcome first; probability average after saturation	Fold probability into severity and silently saturate
EXPECTED_FREQUENCY_OR_EXPOSURE	Expected-input multiplier where incidence/exposure interpretation is justified	Treat as binary event without rationale
ALREADY_EXPECTED_EFFECT	Use $\ell=1$ for the already-integrated probability	Apply probability twice
SCENARIO_CONDITIONAL	Use p_s in the scenario aggregator	Also multiply by ℓ_k
OTHER_DECLARED	Declared formula, units, source and challenge	Outcome-driven operator choice

5.2C Adverse-Confidence Selection Integrity (Normative)

The declared Welfare-stream point estimate MAY retain the pinned confidence-shrinkage construction for continuity and auditability. It is conditional on the declared evidence-support model and MUST NOT by itself authorize unique selection when adverse-confidence sensitivity is unresolved.

An AdverseConfidenceCounterfactual is mandatory for every Tier 3 run, and every high-stakes or decision-material Tier 2 run, when an adverse ordinary-welfare effect has $c_k < 1$ and could contribute to the leading option, Gap, mitigation claim, control-adequacy claim, or option comparison.

Where supported intervals or distributions exist, the counterfactual SHALL use the most adverse supported harm bound and the least favorable supported benefit bound used to justify selection. Where only μ_k and c_k exist, adverse effects use $c_k=1$ in the counterfactual while benefits retain the declared confidence treatment.

This counterfactual is a selection-integrity test, not a claim that every uncertain harm will occur at full magnitude. If the preferred option, Gap classification, or decisive/non-decisive status changes, set CONFIDENCE_RANK_SENSITIVE and prohibit unique framework selection. Public reporting SHALL show both the declared point result and the counterfactual.

A weak adverse estimate **MUST NOT** receive a less protective residual treatment than an honest decision-material unknown solely because c_k is low. Where that would occur, classify the effect as UNKNOWN_IMPACT or use the supported adverse bound and disclose the disposition.

5.3 Temporal Weighting (Normative)

$$\tau(t) = \min(1, \ln(1 + \max(t, t_{\min})) / \ln(1 + T_{\text{ref}}))$$

Intuition anchor (informative): $\tau(t)$ is the fraction of full governance-horizon weight an effect receives.

Defaults:

$t_{\min} = 0.083$ years (approximately 1 month; prevents $\tau(0) = 0$)

$T_{\text{ref}} = 25$ years (governance reference horizon)

Illustrative values (with defaults):

$\tau(100) = 1.00$ (capped)

$\tau(50) = 1.00$ (capped: no bonus beyond T_{ref})

$\tau(25) = 1.00$ (reference horizon: full weight)

$\tau(10) \approx 0.74$ (10 years)

$\tau(5) \approx 0.55$ (5 years)

$\tau(1) \approx 0.21$ (1 year)

$\tau(0.083) \approx 0.025$ (1 month: minimal but nonzero)

Design rationale: Logarithmic weighting preserves intergenerational salience relative to exponential discounting, which aggressively devalues future generations. The cap at T_{ref} prevents the anomaly where distant-future effects receive more weight than the governance horizon itself. T_{ref} is a governed parameter: governance bodies may adjust it (with PCC documentation and sensitivity analysis) but **MUST NOT** set $T_{\text{ref}} < 10$ years without charter-level justification, as this would effectively suppress long-horizon considerations.

Sensitivity note (Tier 3): T_{ref} **SHOULD** be perturbed (for example, $T_{\text{ref}} = 15$ and $T_{\text{ref}} = 50$) as part of the sensitivity bundle. If selection changes under T_{ref} perturbation, the PCC **MUST** document this sensitivity.

Historical change note (v7.5.0; carried forward unchanged in v9.0): This replaces the uncapped $\tau(t) = \ln(1+t)/\ln(1+T_{\text{ref}})$ used in v7.4.5, which produced $\tau(50) \approx 1.21$. The capped version is normatively preferred because temporal weight beyond the governance horizon should not exceed the horizon weight. All prior PCCs computed under v7.4.5 remain valid as lineage material; recomputation is recommended for active decisions with $t > T_{\text{ref}}$ impacts.

The $\min(1, \cdot)$ cap ensures that no effect receives more temporal weight than the reference horizon. Effects at or beyond T_{ref} receive full weight ($\tau = 1.0$); effects below T_{ref} receive proportionally less weight on a logarithmic curve. The $t_{\min}$ floor ensures that very short-term effects are not zeroed out.

Time Horizon (years)	$\tau(t)$ Value
1	≈ 0.21
5	≈ 0.55
10	≈ 0.74
25	1.00

5.3A Short-Duration Peak-Harm Diagnostic (Normative supplement)

Temporal weighting estimates bounded temporal salience and persistence; it is not a cumulative welfare integral. It MUST NOT be interpreted as evidence that an intense short-duration harm is trivial. For every adverse non-rights instance with $|\mu_k| \geq 0.50$ and $\tau(t_k) < 0.10$, the PCC MUST record a PeakHarmReview containing unattenuated peak severity, affected subgroup, reversibility, recurrence risk, mitigation, and whether the harm should enter RF/NCRC, TRC, or CSV instead of remaining ordinary welfare. Tier 3 and high-stakes Tier 2 runs MUST include a sensitivity result that combines the ordinary duration-weighted contribution with a separately disclosed peak-severity diagnostic. This supplement does not convert every brief harm into a rights violation; it prevents duration weighting from hiding intense episodic burdens.

5.3B Temporal Trajectory, Burden, and Token Integrity (Normative Supplement)

The temporal factor $\tau(t)$ is a bounded governance-horizon salience and persistence factor. It is not a measure of cumulative welfare, person-time, recurrence, terminal state, irreversibility, or pure time preference.

A TemporalBurdenReview is required for Tier 3 and high-stakes Tier 2 runs when an effect is recurring, cumulative, persistent beyond T_{ref} , intergenerational, difficult to reverse, terminal-state changing, or materially changes the affected population over time.

The review SHALL record onset, maximal contiguous duration, recurrence, stock-versus-flow status, affected-population trajectory, cumulative exposure where measurable, peak state, terminal state, recovery path, reversibility/irreversibility, intergenerational transmission, uncertainty growth, and any pure-time-preference constitution used.

Temporal token integrity. One continuing state SHALL use its maximal contiguous decision-relevant window. Sequential re-tokenization of one continuous effect into shorter tokens is PROHIBITED unless the episodes are causally or evidentially distinct. A representation-only temporal reversal sets TEMPORAL_PARTITION_SENSITIVE and prevents unique selection.

Temporal pattern	Token treatment	Required evidence / control
Continuous state	One token using the maximal contiguous decision-relevant window	Start, end or review horizon, recovery path, and reopen condition
Recurring exposure without full recovery	One continuing token plus recurrence and cumulative-exposure fields	Recurrence schedule, stock/flow model, and cumulative burden where measurable
Recovered episode followed by a new event	Separate tokens may be used	Evidence of recovery and a distinct causal episode
Distinct causal episodes with separate evidence objects	Separate tokens may be used	Separate pathway, timing, and evidence references
Administrative reporting periods only	View-only subdivisions; no additional RLS mass	Parent token remains controlling

Science may estimate persistence, exposure, recovery, and uncertainty. Any pure preference for present over future welfare is a declared normative constitution and MUST NOT be hidden inside an empirical parameter.

5.4 Pre-Saturation Direct Impact Aggregation (Normative)

For option a, union u , dimension d , define the pre-saturation direct impact per stream $x \in \{\text{base}, \text{welfare}\}$:

$$\tilde{I}_{dir,x}(u,d,a) = \sum_{k \in K(u,d,a)} [r_k \times \tau(t_k) \times \ell_k \times c_k \times e_k \times h_{\{k,x\}} \times \mu_k]$$

where $h_{\{k,\text{base}\}} := 1.0$ for all impact instances. For the Welfare stream, each sensitivity evaluation j uses $h_{\{k,\text{welfare}\}}^{(j)}$ from the governed welfare-inclusion hypothesis set permitted by Section 12 and Appendix

G. MPS bands and intervals inform evidence posture and which hypotheses must be tested; they do not numerically set h_k . Every decision-material hypothesis endpoint MUST be evaluated. When the stream is clear from context, the shorthand $\tilde{I}_{dir}(u,d,a)$ may be used. Appendix B controls the stream-separated equations.

5.5 Saturation (Normative)

To ensure all direct impacts lie in $[-1,+1]$ and avoid runaway totals, RippleLogic uses smooth saturation:

$$I_{dir}(u,d,a) = \tanh(\beta \times \tilde{I}_{dir}(u,d,a))$$

Default: $\beta = 2$.

Saturation-parameter governance (Normative). β is a governed shaping parameter, not a natural constant. Every Tier 2-3 PCC MUST record `beta_direct`; every Tier 3 and high-stakes Tier 2 run MUST report sensitivity over a preregistered range that includes the identity-near case and the canonical default (normally at least $\beta \in \{1.0, 1.5, 2.0, 2.5\}$ unless the domain profile justifies another range). A non-default value MUST be set before outcome inspection, justified by a named domain or governance source, applied symmetrically across compared options, and accompanied by the default counterfactual. If admissibility, ranking, or decisiveness changes materially across the required range, the result is saturation-sensitive and MUST be narrowed, escalated, routed to accountable authority, or treated as non-decisive for the stronger claim.

5.5A Representation and Partition Integrity (Normative)

A representation-preserving change is a split, merge, relabeling, scope allocation, subgroup partition, time slicing, or scenario subdivision that leaves the underlying bearer, state, baseline, causal pathway, evidence object, and total conserved consequence unchanged.

A representation-preserving change MUST NOT be used to increase beneficial RLS contribution or reduce adverse RLS contribution. Within-cell splitting is invariant only when it preserves the same total pre-saturation contribution, probability ownership, evidence and uncertainty semantics, and token identity. Splitting a linear contribution factor can satisfy this condition with all other factors held fixed; splitting a continuing time window generally does not. For example, $\tau(10)=0.736$ but $2\tau(5)=1.100$ and $10\tau(1)=2.127$ under the default temporal rule. Section 5.3B prohibits representation-only temporal re-tokenization. Cross-cell, cross-time, cross-scope, and alternative-tokenization changes require governed review because cell-level nonlinear saturation can otherwise change results.

PartitionInvarianceTest is required for Tier 3 and high-stakes or decision-material Tier 2 runs when: more than one defensible tokenization exists; `BOUNDARY_CONTESTED` is active; a conserved token appears in multiple aggregating cells; a challenger identifies a plausible split/merge; or time/subgroup/scenario granularity can change the number of independently saturated contributions.

These trigger conditions apply to the initial submitted representation as well as to subsequent revisions; the absence of an earlier recorded representation does not establish that the partition test is untriggered.

The test SHALL compare the declared construction with at least one evidence-supported representation-preserving challenger while holding facts, option set, evidence, weights, baseline, and governing parameters constant. If the leading option or decisive/non-decisive status changes, set `PARTITION_RANK_SENSITIVE`, retain all point results for disclosure, and return `REFUSE_DETERMINISTIC_SELECTION` unless the representation is resolved independently of the desired result.

Cell-level tanh remains part of RLS. This section does not authorize moving saturation outside the whole 49-cell field or treating the cells as a natural common welfare unit.

5.6 Magnitude Construction μ_k : Canonical Anchoring Methods (Normative)

Magnitude μ_k MUST represent a baseline-relative change on the normalized $[-1, +1]$ scale. RippleLogic permits two canonical anchoring families:

5.6.1 Percentile Anchoring (Default for non-rights cells)

Let x be the raw indicator (higher-is-better unless specified). Let P_5 , P_{50} , P_{95} be the 5th, 50th, and 95th percentiles of x in a declared reference class.

Define a bounded level score:

If $x \geq P_{50}$:

$$S(x) = \text{clip}((x - P_{50}) / (P_{95} - P_{50}), 0, +1)$$

If $x < P_{50}$:

$$S(x) = \text{clip}((x - P_{50}) / (P_{50} - P_5), -1, 0)$$

Edge-case guard (Normative). If $(P_{95} - P_{50}) = 0$ OR $(P_{50} - P_5) = 0$ (or numerically indistinguishable), percentile anchoring is undefined. The analyst MUST either (a) change the declared reference class, (b) use an alternate canonical anchoring method with explicit PCC justification, or (c) mark the affected impact instances as UNKNOWN_IMPACT and treat the cell as phantom-active under the Phantom Instance Rule until repaired. A run MUST NOT silently divide by zero or substitute arbitrary epsilons without PCC disclosure.

Then define magnitude as a baseline-relative delta:

$$\mu_k = \text{clip}(S(x_a) - S(x_o), -1, +1)$$

Where x_o is baseline indicator value, and x_a is the predicted value under option a.

If higher values are worse (for example, mortality rate), use:

$$S_{\text{worse}}(x) = -S(x)$$

Normative requirement: The PCC MUST record the reference class and percentile values used.

5.6.2 Threshold Anchoring (Required for rights-covered cells unless invariant reference is declared)

For rights-covered cells (cells in any rights coverage set C_r), analysts MUST NOT use context-local percentile anchoring unless the reference class is declared invariant under governance.

Default rule (Tier 1-3): Rights-covered magnitudes MUST be derived using threshold anchoring with explicit “good/bad” anchors so that rights semantics do not drift.

Let x_{good} be the indicator level consistent with rights-safe conditions and x_{bad} be the indicator level representing a severe rights violation onset. Map to a bounded level score:

$$\text{For higher-is-better indicators: } S(x) = \text{clip}((x - x_{\text{bad}}) / (x_{\text{good}} - x_{\text{bad}}), -1, +1)$$

Edge-case guard (Normative). If $(x_{\text{good}} - x_{\text{bad}}) = 0$ (or anchors are otherwise degenerate), threshold anchoring is undefined. The analyst MUST either (a) repair anchors via governed reference conditions, or (b) mark the affected impacts as UNKNOWN_IMPACT and treat the cell as phantom-active under the Phantom Instance Rule until repaired. A run MUST NOT compute a finite score from degenerate anchors.

$$\text{For higher-is-worse indicators: } S(x) = \text{clip}((x_{\text{bad}} - x) / (x_{\text{bad}} - x_{\text{good}}), -1, +1)$$

Then compute baseline-delta magnitude:

$$\mu_k = \text{clip}(S(x_a) - S(x_o), -1, +1)$$

Normative requirement: The PCC MUST record x_{good} , x_{bad} , indicator definition, and direction.

Magnitude construction for rights-covered cells (C_r cells). For rights-covered cells, magnitudes **MUST** be derived using threshold anchoring with explicit reference conditions aligned to rights semantics (Section 7.1.1), preventing the context-local drift that percentile anchoring would permit. The PCC **MUST** record x_{good} , x_{bad} , indicator definition, direction, and the source/justification for the anchor values. This requirement applies to all tiers when rights-covered cells are quantitatively assessed.

5.7 Missing Data Rule (Ignorance Penalty) (Normative)

To prevent score inflation by omission:

If a welfare cell (u,d) is active (required by tier context and not masked for RLS), but no defensible instances can be specified, the PCC **MUST**:

- Mark the cell as “UNKNOWN_IMPACT”, and
- Include a phantom instance with canonical parameters:

Parameter	Phantom Value
μ_{phantom}	-0.10
r	1
t	T_{ref} (25 years)
ℓ	1
c	1
e	1
$h_{\{k,\text{base}\}}$	1
$h_{\{k,\text{welfare}\}}$	1

This yields a mild negative default rather than unjustified neutrality.

Calibration rationale (Informative). The phantom magnitude $\mu_{\text{phantom}} = -0.10$ is a governance prior set to satisfy two constraints simultaneously: it must be negative enough that omitting evidence is worse than reporting neutral impact (preventing score inflation by ignorance), and it must be small enough in absolute value that phantom instances do not dominate well-evidenced cells when multiple unknowns coexist in a run. At $\mu = -0.10$ with default parameters ($r=1$, $t=T_{\text{ref}}$, $\ell=1$, $c=1$), the phantom produces a saturated direct impact of approximately -0.20 per cell. In a run with uniform weights across 49 cells, a single phantom cell contributes roughly 0.004 to the total RLS magnitude, which is detectable in audit but not decisive unless many cells are unknown simultaneously. If roughly five to seven cells carry phantom instances under uniform 49-cell weights, the cumulative penalty is approximately -0.02 to -0.03 in RLS; additional phantom-active cells increase the magnitude further. This creates a meaningful incentive to gather evidence without treating the convention as empirical harm evidence. The value -0.10 is a governance prior, not an empirically validated optimum. It is explicitly challengeable under standard governance rules; any revision requires justification, version increment, and sensitivity analysis over a reference decision suite.

Phantom-dependence rule (Normative). A phantom value is an ignorance-management convention, not evidence that harm occurred. For Tier 3 and decision-material or high-stakes Tier 2 runs, the PCC **MUST** report the number and locations of phantom-active cells, the RLS ranking with and without phantom instances, and a governed unknown or adverse-bound sensitivity where feasible. If the preferred option, decisiveness classification, or public conclusion changes materially because of the phantom convention, the result **MUST** be treated as phantom-sensitive and therefore non-decisive, narrowed, escalated, or refused for the stronger claim. Analysts **MUST NOT** add weak estimates merely to avoid the phantom rule or manipulate the active-cell set to advantage an option.

Clarification: This rule does not apply to cells that are legitimately out of scope and properly masked with justification; it DOES apply to non-maskable cells and to any cell required by the tier's minimum coverage rules.

Gate-Critical Unknown Rule (Normative). If a cell within C_r or C_{cat} is declared UNKNOWN_IMPACT, admissibility for that cell SHALL NOT be resolved using only the standard phantom-instance fallback. Default handling path: for rights-covered cells, the run MUST either (i) apply a governed conservative bound sufficient to preserve the rights gate, or (ii) treat the gate as unresolved and rerun, downgrade, or escalate. For catastrophe-covered cells, the run MUST either (i) apply a governed conservative scenario/cell bound sufficient to preserve TRC, or (ii) treat TRC as unresolved and rerun, downgrade, or escalate. Implementations MUST set RIGHTS_CELL_UNKNOWN_GATE_CRITICAL or CATASTROPHE_CELL_UNKNOWN_GATE_CRITICAL, as applicable, using only canonical severities defined in Section 14.3.

5.8 Scenario-Conditioned Impacts (Normative for TRC runs)

When TRC is in use (Tier 3 required; Tier 2 required when catastrophe relevance plausible), impacts MUST be scenario-conditioned at least for catastrophe cells:

Scenario probabilities p_s are governed (Appendix D). Scenario conditioning enters through ℓ_k and/or instance presence, and through propagation if scenario-specific kernels are declared.

5.9 Worked Examples (Non-Normative, Computation-Illustrative)

Example 5.1: Direct impact, Community-Social cell

Decision: Remote-work policy option a. Cell: $u=3$ (Community), $d=3$ (Social).

Assume $T_{ref} = 25$, $\beta = 2$. Here likelihood terms denote expected frequency/exposure, not an all-or-nothing event probability. Use unrounded values in computation.

Instance 1: Reduced in-person interaction

- $\mu_1 = -0.30$, $r_1 = 0.80$, $t_1 = 3$ years, $\ell_1 = 0.90$, $c_1 = 0.70$, $e_1 = 1$, $h_{\{1,base\}} = 1$, $h_{\{1,welfare\}} = 1$

Instance 2: Increased online community participation

- $\mu_2 = +0.15$, $r_2 = 0.50$, $t_2 = 5$ years, $\ell_2 = 0.70$, $c_2 = 0.50$, $e_2 = 1$, $h_{\{2,base\}} = 1$, $h_{\{2,welfare\}} = 1$

Compute temporal weights:

- $\tau(3) = \ln(4)/\ln(26) \approx 0.43$
- $\tau(5) = \ln(6)/\ln(26) \approx 0.55$

Compute contributions:

- Inst1: $0.80 \times [\ln(4)/\ln(26)] \times 0.90 \times 0.70 \times 1 \times 1 \times (-0.30) \approx -0.064334$
- Inst2: $0.50 \times [\ln(6)/\ln(26)] \times 0.70 \times 0.50 \times 1 \times 1 \times (+0.15) \approx +0.014436$

Aggregate pre-saturation:

- $\tilde{I}_{dir} \approx -0.064334 + 0.014436 \approx -0.049898$

Saturate:

- $I_{dir}(3,3,a) \approx \tanh(2 \times -0.0498984666) \approx -0.099467$

Example 5.2: Missing-data penalty

Cell $u=7$ (Biosphere), $d=7$ (Environment) is active (non-maskable for environment-relevant decisions). Analyst lacks defensible estimates.

Phantom instance yields:

- $\tilde{I}_{dir} = 1 \times 1 \times 1 \times 1 \times 1 \times 1 \times (-0.10) = -0.10$
- $I_{dir} = \tanh(2 \times -0.10) \approx -0.20$

This is a declared ignorance-penalty display, not evidence of actual harm or a gate PASS; UNKNOWN_IMPACT and the required uncertainty/review handling remain controlling.

SECTION 6: RIPPLE PROPAGATION: KERNEL AND EPISTEMIC HUMILITY

6.0 Purpose (Normative)

Ripple propagation models how direct impacts in one cell causally and institutionally ripple into other cells across unions and dimensions. This makes cross-scale externalities explicit and auditable.

Tier posture: Tier 1-3 permit NONE or QUICK propagation. FULL propagation is a Tier-4 design target only and MUST NOT be used for Tier 1-3 compliance claims.

6.1 Kernel Definition and Convention (Normative)

Let the 49 welfare cells be indexed by $i = \varphi(u,d)$ with $\varphi(u,d) = 7(u-1) + d$. The ripple kernel is a sparse matrix $K \in \mathbb{R}^{49 \times 49}$.

Kernel convention (MUST):

- $K_{\{ij\}}$ maps effect from source cell j to target cell i (target-row, source-column).
- Propagation uses left multiplication: $\tilde{I}_{\text{prop}} = I_{\text{dir}} + K \times I_{\text{dir}}$

Interpretation:

Scale semantics and first-order propagation note (Normative clarification). The bounded cell impacts used by RippleLogic are governance scores anchored to declared indicators and baselines, not claims that dignity, liberty, or meaning are measured as cardinal utility in nature. The kernel K is a declared first-order propagation operator for bounded local approximation under the chosen decision horizon. It does not claim that all socio-technical dynamics are globally linear. Non-linearities and threshold behavior are handled elsewhere in the architecture through saturation, scenario branching, rights floors, tail-risk screening, containment, escalation, and versioned recalibration of K . Implementations MUST NOT present K as a proof of empirical linearity; they MAY present it as a transparent, auditable approximation layer governed by the PCC.

- $K_{ij} > 0$: improving source j tends to improve target i .
- $K_{ij} < 0$: improving source j tends to harm target i .
- $K_{ij} = 0$: no modeled pathway.

6.1A Scope Blocks, Causal Provenance and No Duplicate Propagation (Normative)

Use the existing flattening map $\varphi(u,d)=7(u-1)+d$. The 49 by 49 kernel can be read as a seven by seven array of seven by seven blocks: $K^{(u \leftarrow v)}[d,d_{\text{prime}}] := K[\varphi(u,d),\varphi(v,d_{\text{prime}})]$. Target rows and source columns are unchanged. $T_{(v \rightarrow u)}$, where used in an explanatory diagram, is only an alias for this block, not a second propagation operator. Cross-scope blocks are not identity maps by default; an unmodelled edge is not evidence that no pathway exists.

Each relied-upon edge SHALL reference its source and target constructs, effect/pathway evidence, causal direction and sign, coefficient and units or normalization, lag and horizon, configuration/domain, uncertainty, validity and review trigger. This is the provenance extension of existing KQS/KOPS and PCC records, not a replacement kernel-quality gate. Do not infer a coefficient from visual taxonomy, raw compute load, blockchain consensus or a linguistic analogy.

Every direct-versus-propagated construction SHALL state whether a downstream effect is already included. A direct total-effect estimate and K times its upstream cause cannot both add the same endpoint consequence.

Shared causes, mediation, feedback and repeated observations require explicit attribution or exclusion rules. This rule preserves genuine distinct downstream outcomes rather than treating causality itself as duplication.

A state-dependent susceptibility modifier, if researched, is structural telemetry rather than a Welfare Dimension, sentience estimate or hidden e_k /weight multiplier. Any numerical use requires a separately warranted model, fixed mapping, joint uncertainty and tests of the final effective kernel against the existing stability, KQS and sensitivity conditions. No Gamma vector, inverse reserve formula or automatic cross-substrate welfare mapping is adopted by this specification.

Diagonal-edge clarification. A nonzero K_{ji} is not automatically a new effect or additional welfare entitlement. It requires a distinct, evidence-supported within-cell secondary pathway and lag, the ordinary kernel bounds, and an explicit check that the endpoint is not already included in direct impacts. With no such warrant, use $K_{ji}=0$. This is the existing causal-provenance and no-duplicate-propagation rule applied to the diagonal, not a new propagation mode.

6.2 Propagation Modes (Tier 1-3) (Normative)

Mode	Formula	Use Case
NONE (Direct-only)	$I_{prop} := I_{dir}$	Tier 1-2 default
QUICK (First-order)	$\tilde{I}_{prop} = I_{dir} + K \times I_{dir}$	Tier 3 with sensitivity

Stream note (Normative). Propagation is applied per stream $x \in \{\text{base}, \text{welfare}\}$ as defined in Appendix B. In Section 6 (and any tables using shorthand), I_{dir} and I_{prop} are shorthand for $I_{dir,x}$ and $I_{prop,x}$ when the stream is clear from context. Admissibility gates (NCRC, TRC, CSV) use the Base stream; RLS uses the Welfare stream.

Normative restriction: FULL propagation is PROHIBITED for Tier 1-3 claims in the current v13.0 release line.

6.3 Post-Propagation Saturation (Normative)

Because propagation can push values outside $[-1, +1]$, apply elementwise saturation:

$$I_{prop_x}(u,d,a) = \tanh(\beta_{prop} \times \tilde{I}_{prop_x}(u,d,a))$$

Notation: I_{prop_base} and $I_{prop_welfare}$ denote the $x=\text{base}$ and $x=\text{welfare}$ outputs respectively.

Default: $\beta_{prop} = 1$.

Propagation-mode comparability rule (Normative clarification). NONE and QUICK are different declared transformations: QUICK includes a second saturation step while NONE preserves the direct saturated value. Consequently, a NONE-versus-QUICK comparison is not a pure estimate of the kernel term alone, and a humility fallback from QUICK to NONE is a model change rather than simple deletion of K . The PCC MUST disclose that distinction, recompute the complete result under each declared mode, and separate kernel sensitivity from mode/saturation sensitivity. If the preferred option or gate interpretation changes because of the mode boundary, the stronger deterministic claim MUST be narrowed or treated as unresolved. A future formula change intended to make the null-kernel limit identical across modes requires a versioned equation change and reference-suite replay; it MUST NOT be introduced silently as a same-run implementation convenience.

6.4 Stability Constraints and Humility Fallbacks (Normative)

Kernel stability guardrails (defaults; may be tightened by governance):

- Entry bound: $|K_{\{ij\}}| \leq k_{max}$, default $k_{max} = 0.5$
- Absolute row-sum bound (required): $\|K\|_{\infty} := \max_i \sum_j |K_{\{ij\}}| \leq \rho_{max}$, default $\rho_{max} = 0.9$.

- Spectral radius constraint (satisfied by construction under the row-sum bound): If $\|K\|_\infty < 1$, then $\rho(K) < 1$ is treated as satisfied for Tier 1-3 purposes because $\rho(K) \leq \|K\|_\infty$. Implementations MUST NOT require additional spectral-radius computation to claim Tier 1-3 conformance when the row-sum bound holds.

Kernel stability verification rule (Normative). A Tier 1-3 implementation MUST compute and record (in PCC when kernel is used): (i) $\max_{ij} |K_{ij}|$ and (ii) $\|K\|_\infty$. If either exceeds its bound (entry bound or row-sum bound), the run MUST execute the humility fallback (`propagation_mode = NONE`) and record `KERNEL_HUMILITY_FALLBACK = TRUE`.

If the declared kernel violates entry bounds or row-sum bounds (and therefore fails the Tier 1-3 stability verification rule above), the run MUST:

- Set `propagation_mode = NONE`, and
- Record a PCC limitation: `KERNEL_HUMILITY_FALLBACK = TRUE`.

Rationale: An unstable or over-amplifying kernel is worse than no kernel; it creates false certainty and is easy to game.

6.5 Kernel Quality Score (KQS) (Normative)

KQS summarizes readiness of the kernel for decision-relevant propagation:

$$\text{KQS} = w_{\text{cov}} \times C_{\text{cov}} + w_{\text{id}} \times C_{\text{id}} + w_{\text{stab}} \times C_{\text{stab}} + w_{\text{pred}} \times C_{\text{pred}}$$

Default component weights:

Component	Weight
Coverage (C_{cov})	0.25
Identifiability (C_{id})	0.30
Stability (C_{stab})	0.20
Prediction (C_{pred})	0.25

Component meanings (each in $[0,1]$):

- C_{cov} : coverage evidence proportion (share of relied-upon edges with cited evidence)
- C_{id} : identifiability (edges are specified with clear endpoints and sign; replayable)
- C_{stab} : stability margin (satisfies bounds with margin)
- C_{pred} : predictive accuracy (backtest/pilot performance where available; otherwise conservative prior)

6.5.1 Canonical KQS computation (Tier 1-3 deterministic) (Normative)

Definition (Normative). A ‘relied-upon edge’ is any kernel edge (i,j) with $|K_{ij}| \geq k_{\text{use_min}}$ that touches at least one active cell (Section 10.3) or any non-maskable cell in C_r or C_{cat} . Default $k_{\text{use_min}} = 0.05$ unless the PCC declares a stricter value.

C_{cov} (Normative). Let E_{use} be the set of relied-upon edges. Let E_{evid} be those edges in E_{use} whose evidence fields include (i) an `EvidenceClass` $\in \{A,B,C\}$ and (ii) a citation or provenance reference sufficient for retrieval. Then $C_{\text{cov}} := |E_{\text{evid}}| / \max(1, |E_{\text{use}}|)$.

C_{id} (Normative). Let E_{id} be those edges in E_{use} with declared source union/dimension, target union/dimension, sign, and numeric magnitude fields present and within canonical bounds. Then $C_{\text{id}} := |E_{\text{id}}| / \max(1, |E_{\text{use}}|)$.

C_{stab} (Normative). If QUICK is used, compute `PCC.Kernel.RowSumAbsMax` = $\|K\|_\infty$ and `PCC.Kernel.EntryAbsMax` = $\max_{ij} |K_{ij}|$ per Section 6.4. Set $C_{\text{stab}} := 0$ if any bound is violated; otherwise $C_{\text{stab}} := \text{clip}((\rho_{\text{max}} - \|K\|_\infty) / \rho_{\text{max}}, 0, 1)$, with default $\rho_{\text{max}} = 0.9$.

C_pred (Normative). If no backtest or pilot evidence exists for the decision domain, set $C_pred := 0.30$ (NO_EVIDENCE_DEFAULT) as already required. If empirical evidence exists, the PCC MUST declare the predictive metric used and compute C_pred as $\text{clip}((\text{MetricScore} - \text{MetricBaseline}) / (\text{MetricTarget} - \text{MetricBaseline}), 0, 1)$ with all three values disclosed in PCC so an auditor can recompute.

Disclosure rule (Normative). The PCC MUST record k_use_min , the edge counts $|E_use|$, $|E_evid|$, $|E_id|$, and the computed component values C_cov , C_id , C_stab , C_pred used in the KQS calculation.

Default C_pred prior (Normative). If no backtest/pilot evidence exists for the decision domain, set $C_pred := 0.30$ (NO_EVIDENCE_DEFAULT). The PCC MUST (i) mark this basis as NO_EVIDENCE_DEFAULT, and (ii) include a plan and timeline to obtain empirical validation evidence; until such evidence exists, C_pred SHALL NOT be reported as higher than the prior.

Interpretation of KQS components (Normative clarification). C_id measures the specified field completeness and replayability, not statistical or causal identification. $C_pred = \text{NO_EVIDENCE_DEFAULT}$ is a declared prior, not measured predictive accuracy. C_stab measures margin below the stated numerical bound, not empirical validity of a causal model. Aggregate KQS eligibility does not establish predictive validation or override Reality Grounding, uncertainty, evidence-domain or robustness requirements. The existing component equations, eligibility rules and sensitivity requirements remain unchanged.

KQS policy (MUST) for Tier 3:

- **Evidence-coverage eligibility floor:** if $C_cov < 0.40$, the kernel MUST NOT be used for QUICK propagation regardless of the aggregate KQS; set `propagation_mode = NONE` and record `KERNEL_EVIDENCE_COVERAGE_INSUFFICIENT`. Formal field completeness, stability, or a no-evidence prediction prior cannot compensate for inadequate evidence coverage.
- If $KQS < 0.40$: Kernel MUST NOT be used; `propagation_mode = NONE`.
- If $0.40 \leq KQS < 0.50$: QUICK allowed only with mandatory kernel sensitivity; otherwise use NONE.
- If $KQS \geq 0.50$: Kernel use permitted with required sensitivity at Tier 3.

Tier 1-2 default: `propagation_mode = NONE` unless the PCC explicitly declares kernel use and its KQS.

6.6 Kernel Sensitivity Requirements (Normative for Tier 3 when QUICK is used)

When QUICK propagation is used at Tier 3, the PCC MUST include sensitivity analysis:

- Perturb each relied-upon edge in E_use (as defined in Section 6.5.1) by ± 0.05 (or $\pm 10\%$ of its magnitude, whichever is larger; declare the rule).
- Recompute the cascade for each perturbation set (at minimum one-at-a-time).
- If admissibility outcomes (NCRC, TRC, CSV) or the selected option changes, set `audit_flag = DECISION_FRAGILE_KERNEL` and escalate per tier policy.

6.7 KOPS and Starter Kernel (Tier 2-3) (Normative plus Provisional Labeling)

RippleLogic permits a governed subset of kernel edges called the Key Operational Pathways Set (KOPS). KOPS is the set of “load-bearing” pathways that are documented with evidence notes, sign-checked against literature or backtests, and sensitivity-audited.

A Starter KOPS may be provided (Appendix K) and MUST be labeled PROVISIONAL with evidence classes and a global shrink factor for elicited edges (default 0.35). If a starter kernel is used, the PCC MUST disclose this and MUST include kernel sensitivity.

6.7A Continuous-Time Research Bridge (Informative; not a new propagation mode)

A separately warranted physical or institutional state model may use $dz/dt = M_dyn z + B_dyn a + xi$. For a locally linear, time-invariant unforced model, $z(t + \Delta t) = \exp(M_dyn \Delta t) z(t)$, with M_dyn carrying inverse-time

units in consistently normalized coordinates. The approximation $\exp(M_{\text{dyn}} \Delta t)$ approximately equals $I_d + M_{\text{dyn}} \Delta t$ is local in $M_{\text{dyn}} \Delta t$ and does not justify universal linear welfare dynamics. The existing K may be compared with $M_{\text{dyn}} \Delta t$ only after the state-to-impact bridge, units, horizon, saturation and omission treatment have been reconciled.

This research relation does not replace NONE or QUICK, certify a matrix exponential as stable, or make FULL/Tier-4 claims available. Positive feedback can amplify losses; recovery must be modeled as an intervention with resource costs, delays, failure modes and externalities. Neither a harmonic-amplification theorem nor reversal of physical entropy is asserted.

6.8 Scenario-Conditioned Propagation (Tier 1-3 Default Rule)

Default: Kernel is not scenario-conditioned; scenario enters via instance likelihoods ℓ_k and scenario-specific instance activation. Scenario-conditioned kernels K_s may be used only with explicit PCC justification and must carry KQS and sensitivity per scenario class.

SECTION 7: RIGHTS LAYER: NON-COMPENSATORY RIGHTS CONSTRAINT (NCRC)

7.0 Purpose and Union-Based Ethics Rationale (Normative)

RippleLogic enforces non-compensability: certain harms are not permitted to be “paid for” by aggregate benefits elsewhere. NCRC operationalizes this by removing rights-violating options prior to any welfare optimization.

Union-Based Ethics (UBE) justification: Under UBR, harms propagate across layered and intersecting analytical scopes and dependency paths; without a rights-first constraint, local optimization systematically externalizes costs onto vulnerable subgroups and future stakeholders. NCRC is therefore a structural safeguard against value scalarization and moral laundering (Rawls, 1971; Sen, 2009).

7.1 Canonical Rights Set (Normative)

RippleLogic defines eight core rights as feasibility floors:

Table 7.1: Canonical Rights Thresholds (Normative)

Right	Code	Threshold θ_r	Normative Anchor
Life	LIFE	-0.90	Near-certain or highly probable death or lethal exposure
Bodily Integrity	BODY	-0.70	Severe injury, disability, torture, serious bodily violation
Liberty	LBTY	-0.65	Arbitrary detention, forced labor, coercive confinement
Basic Needs	NEED	-0.50	Severe deprivation: food insecurity, homelessness, loss of basic subsistence
Dignity	DIGN	-0.55	Systematic humiliation, dehumanization, targeted degradation
Due Process	PROC	-0.45	Denial of fair hearing, non-transparent coercive procedure
Information	INFO	-0.40	Systematic censorship, epistemic coercion, pervasive disinformation constraints
Ecological Integrity	ECOL	-0.65	Material biosphere integrity breach or planetary boundary transgression corridor

Ecological Integrity boundary. ECOL is a MathGov non-compensatory enabling-condition protection for biosphere integrity and life-supporting ecological conditions. It may overlap with legal environmental rights where recognized, but this Canon does not claim universal legal status by itself.

Normative interpretation rule: Thresholds θ_r are admissibility floors, not weights. They do not represent “how important” a right is; they represent the minimum allowable protection level.

Threshold governance. All thresholds are explicitly challengeable. Revisions require justification, version increment, PCC documentation preserving prior values, and sensitivity outcomes over a reference decision suite.

Threshold epistemic boundary. Rights thresholds are versioned governance priors, not discovered empirical constants or jurisdiction-specific legal conclusions. They may be challenged, localized, stress-tested, and revised under governed update rules, but they remain non-compensatory while active in a run.

Sensitivity requirement. In Tier 3 runs, each θ_r MUST be perturbed by ± 0.05 and the cascade recomputed. If admissibility changes for any option under perturbation, the decision is flagged as threshold-sensitive and the PCC MUST document which rights are near-binding, which options are affected, and how the final selection responds to that near-binding status.

Proxy disclaimer. Governance indices (such as press freedom or rule-of-law indices) may be used as operational proxies for anchoring reference conditions, but are not treated as moral ground truth and MUST be accompanied by documented limitations in the PCC.

Right INFO ($\theta = -0.40$). Normative anchor: systematic censorship, epistemic coercion, or deliberate degradation of the information environment such that agents cannot form informed preferences. Example anchors: systematic media control and censorship per RSF Press Freedom Index “Very Serious Situation” (score below 40); persistent large-scale disinformation operations affecting public health, safety, or governance; surveillance-driven chilling effects that measurably suppress information seeking or expression.

Right PROC ($\theta = -0.45$). Normative anchor: denial of fair hearing, non-transparent coercive procedure, or absence of meaningful recourse against institutional power. Example anchors: no meaningful access to justice or independent review per WJP Rule of Law Index bottom quintile; consequential automated decisions without explanation, appeal, or human review; systematic procedural exclusion; coercive institutional action without proportionality review.

Right NEED ($\theta = -0.50$). Normative anchor: severe deprivation of material conditions necessary for survival and minimally adequate functioning. Example anchors: sustained severe food insecurity at IPC Phase 3 or above; homelessness or severe housing deprivation per UN-Habitat standards; lack of safe water and sanitation; denial of essential healthcare access such that treatable conditions become severe or lethal.

Right DIGN ($\theta = -0.55$). Normative anchor: systematic humiliation, dehumanization, or targeted degradation that denies inherent worth and equal standing. Example anchors: institutionalized dehumanization of groups per the UN Framework of Analysis for Atrocity Crimes; degrading treatment per ECHR Article 3; systematic exclusion from participation in social and civic life; discriminatory institutions that deny equal standing below capability thresholds (Nussbaum, 2011).

Right ECOL ($\theta = -0.65$). Normative anchor: breach of biosphere integrity or entry into high-risk Earth-system transgression corridors that threaten enabling conditions for continued life. Example anchors: actions that materially increase probability of crossing critical climate or ecosystem tipping elements per Lenton et al. (2008); ecosystem regime shifts effectively irreversible on human-relevant timescales; depletion of foundational natural capital (aquifers, soils, fisheries) beyond recharge capacity; biodiversity loss sufficient to undermine ecosystem services.

Right LBTY ($\theta = -0.65$). Normative anchor: arbitrary detention, forced labor, coercive confinement, or systematic deprivation of freedom of movement and association. Example anchors: deprivation of liberty without legal basis or review per ICCPR Article 9; forced labor patterns consistent with ILO indicator bundles

(restriction of movement, debt bondage, retention of identity documents, isolation); coercive confinement not justified by legitimate, proportionate public safety measures.

Right BODY ($\theta = -0.70$). Normative anchor: severe injury, disability, torture, or serious violation of bodily autonomy falling short of lethality but constituting fundamental harm to physical and psychological functioning. Example anchors: torture or cruel treatment as defined by the UN Convention Against Torture; serious bodily injury resulting in permanent disability or substantial loss of function per WHO disability classification frameworks (ICF; WHODAS 2.0); forced medical procedures or systematic bodily violations; sustained exposure to hazards expected to cause serious chronic illness.

Right LIFE ($\theta = -0.90$). Normative anchor: near-certain or highly probable death, lethal exposure, or elimination of conditions necessary for continued biological existence for the worst-off subgroup. Example anchors: direct lethal violence or exposure; catastrophic deprivation conditions such as IPC Phase 5 famine ($\geq 20\%$ extreme food gaps, crude death rate >2 per 10,000 per day); collapse of life-sustaining infrastructure (water, sanitation, shelter) without viable alternatives; mass-casualty disaster exposures. The threshold is set above -1.0 because -1.0 corresponds to guaranteed death for all affected persons, a theoretical maximum rarely achieved even in extreme scenarios.

7.1.1 Threshold Calibration Protocol (Normative)

The eight canonical rights thresholds are normative governance priors, not empirically derived frequencies. Each threshold encodes a severity level on the normalized impact scale $[-1, +1]$ at which a predicted impact is treated as a rights-floor violation that cannot be traded away for aggregate benefit elsewhere.

Calibration follows two principles:

First, thresholds are placed relative to the semantics of the impact scale: -1.0 denotes an extreme, worst-case welfare decrement in the relevant domain; $+1.0$ denotes an extreme improvement.

Second, the ordering of thresholds reflects both a severity gradient and an enabling-conditions logic: rights protecting against lethal and bodily harms sit closer to the floor, while rights protecting procedural and epistemic conditions are set closer to zero because their violation can enable cascading downstream rights failures. When information environments are degraded or due process is absent, communities lose the epistemic and procedural capacity to detect and contest violations of life, bodily integrity, and liberty.

Operationally, a threshold is justified by two to four anchor conditions per right, a mapping rule from anchors to the $[-1, +1]$ scale using threshold anchoring (Section 5.6.2), and a mandatory sensitivity sweep.

7.2 Rights Coverage Sets C_r (Normative)

Coverage interpretation. Absence of D6 from the default floor-coverage sets is not an exemption from rights review. A Meaning-dimension finding that evidences a categorical, severe-hazard or other protected-rights concern reopens RF/NCRC under the applicable rights profile and WDBIP routing rules; retain the effect identity and do not duplicate welfare mass merely to create a rights view.

Each right applies to a defined subset of welfare cells C_r . NCRC evaluates each right across its coverage set and uses the worst-off subgroup for rights-covered cells.

Canonical coverage sets (authoritative):

Coverage note (Normative clarification). ECOL is scoped to capture enabling-conditions harms to the biosphere and to humanity's life-support systems. Canonically this is implemented via Environment-dimension cells in the Humanity / Global Coordination and Biosphere unions (and any additional declared cells via governed extension). Implementations MAY extend ECOL coverage to additional cells only by explicit governance declaration in the PCC, and MUST NOT silently re-scope ECOL between runs.

Right	Coverage Set C_r
LIFE	$\{(u, \text{Health}): u \in \{1,2,3,4,5,6\}\} \cup \{(6, \text{Environment})\}$

Right	Coverage Set C _r
BODY	$\{(u, \text{Health}): u \in \{1,2,3,4,5,6\}\}$
LBTY	$\{(u, \text{Agency}): u \in \{1,2,3,4,5,6\}\} \cup \{(u, \text{Social}): u \in \{3,4,5,6\}\}$
NEED	$\{(u, \text{Material}): u \in \{1,2,3,4,5,6\}\} \cup \{(u, \text{Health}): u \in \{1,2,3,4,5,6\}\}$
DIGN	$\{(u, \text{Social}): u \in \{1,2,3,4,5,6\}\} \cup \{(u, \text{Agency}): u \in \{1,2,3,4,5,6\}\}$
PROC	$\{(u, \text{Agency}): u \in \{4,5,6\}\} \cup \{(u, \text{Knowledge}): u \in \{4,5,6\}\} \cup \{(u, \text{Social}): u \in \{4,5,6\}\}$
INFO	$\{(u, \text{Knowledge}): u \in \{1,2,3,4,5,6\}\} \cup \{(u, \text{Agency}): u \in \{1,2,3,4,5,6\}\}$
ECOL	$\{(6, \text{Environment}), (7, \text{Environment})\}$

7.3 Worst-Off Subgroup Semantics (Normative)

Rights cannot be averaged away. For every rights-covered cell (u,d) in any C_r, define:

$$I_{\text{rights}}(u,d,a,r) = \min_{\{g \in G_{\text{RF}}(u,d,a,r)\}} I_{\text{RF}}(u,d,a|g,r)$$

G_{RF} MUST NOT be empty for a material rights-covered cell. When no governed subgroup split is material, the declared base population is the sole member. If neither a governed subgroup record nor a defensible base-population record exists, the cell is NCRC_UNKNOWN; an empty set MUST NOT pass vacuously.

Where G_{RF}(u,d,a,r) is the governed set of protected subgroups relevant to that right and cell. The right identifier r MUST remain explicit in machine records even where display equations suppress it for readability. The ordinary Welfare/Base-stream value I_{prop_base} may inform evidence discovery, but it is not the canonical rights-floor input unless separately reconstructed under the Section 7.4 non-attenuation rules.

Tier requirements:

- Tier 3: Subgroup enumeration is REQUIRED for rights-covered cells (Appendix H PCC requirements).
- Tier 2 (strict posture): Subgroup enumeration is REQUIRED at least for directly affected populations in rights-covered cells; when infeasible, the PCC MUST record SUBGROUP_LIMITATION and apply a conservative bound rule (Section 7.3.2).

7.3.1 Minimum subgroup categories (Normative minimum)

At minimum, when applicable and feasible, G_{u,d} SHOULD include subgroups defined by:

- Age (children; elderly)
- Disability status
- Legally protected characteristics relevant in jurisdiction
- Economic vulnerability (for example, bottom income quintile)
- Geographic exposure (for example, high-risk locations)
- Domain-specific protected groups (for example, patients, detainees, precarious workers, indigenous communities)

7.3.2 Conservative fallback when subgroup disaggregation is infeasible (Tier 2-3) (Normative)

If subgroup-specific rights-floor impacts I_{RF}(u,d,a|g,r) cannot be produced for a rights-covered cell, the run MUST:

- Record SUBGROUP_LIMITATION in the PCC, and
- Apply a conservative rights bound for that cell:

Let γ_{subgroup} be a conservatism factor (default γ_{subgroup} = 1.5). Then define:

- If the separately constructed base-population rights-floor impact I_{RF_base}(u,d,a,r) < 0, set I_{rights}(u,d,a,r) = max(-1, γ_{subgroup} × I_{RF_base}(u,d,a,r))
- Else set I_{rights}(u,d,a,r) = I_{RF_base}(u,d,a,r)

This bound applies only for NCRC checking, not for RLS scoring.

Fallback applicability (Normative clarification). A nonnegative welfare average does not establish the absence of protected-subgroup harm. SUBGROUP_LIMITATION does not waive material subgroup closure. Where a decision-material subgroup harm remains unresolved and no defensible rights bound can be established, the run MUST use NCRC_UNKNOWN and the existing escalation or refusal route rather than treating the fallback as evidence of passage. The governed γ_{subgroup} factor is a screening convention, not a distribution-free guarantee for every unobserved subgroup. The adverse-only rights construction in Section 7.4 remains controlling.

Emergency comparison tie rule (Normative). The saturated I_{rights} value remains controlling for ordinary admissibility. If two or more Emergency Mode options tie on v_r^{floor} because the conservative clamp maps them to -1 , the run MUST use the unclamped conservative value $\gamma_{\text{subgroup}} \times I_{\text{RF_base}}$ as a secondary ordering key for that right, with more negative values treated as worse and the trace recorded in the PCC. This secondary key does not convert any failing option into an ordinary NCRC pass.

If $I_{\text{RF_base}}$ cannot be constructed from reviewed rights-floor instances, the cell is NCRC_UNKNOWN; the run MUST collect evidence, narrow, escalate, or refuse rather than reuse a temporally or confidence-attenuated welfare value as a rights pass.

Calibration note (Normative guidance). The default $\gamma_{\text{subgroup}} = 1.5$ is a conservative governance prior rather than an empirically estimated constant. The value is set as a moderate worst-off uplift between the identity case ($\gamma_{\text{subgroup}} = 1.0$) and stronger stress-test settings used to examine subgroup vulnerability under uncertainty. It is intended to reduce subgroup erasure without asserting a precise empirical constant across domains. Tier 3 deployments SHOULD review γ_{subgroup} against their reference decision suite at least annually or after a material accumulation of rights-sensitive runs. Any non-default γ_{subgroup} MUST be declared in the PCC, justified, and accompanied by sensitivity reruns over the reference suite. The sensitivity sweep for γ_{subgroup} SHOULD include at minimum $\gamma_{\text{subgroup}} \in \{1.0, 1.3, 1.5, 2.0\}$ to test whether admissibility outcomes or selection change materially.

7.4 Rights-Specific Severity, Hazard, and NCRC Admissibility (Normative)

7.4.1 Rights-floor impact channel (Normative)

The welfare-impact stream is not by itself sufficient to decide rights admissibility. For every adverse impact instance k that is material to a protected right r , the run MUST construct a rights-floor instance using the exposed protected subgroup as the reach denominator.

For adverse rights-floor instances, set the rights temporal factor $\tau_{\text{RF}}(t_k) := 1.0$. Duration remains recorded and may aggravate severity, remedy, and emergency comparison, but short duration MUST NOT attenuate a categorical or severe rights violation into admissibility. Analyst confidence c_k is evidence metadata rather than a protective discount: low confidence MUST trigger conservative bounding, evidence collection, narrowing, escalation, or refusal, not a weaker rights check.

Let ℓ_k^{U} be the declared conservative upper likelihood bound for the rights event. The rights-floor contribution is:

$$x_{\text{RF}}(k, r, a) = \mu_k \cdot r_k^{\text{R}} \cdot \ell_k^{\text{U}} \cdot e_k$$

where r_k^{R} uses rights-specific exposed-subgroup reach. Only adverse rights-floor instances enter this channel; positive effects remain in welfare analysis and cannot offset a protected loss.

For each rights-covered cell, let $K_{\text{RF}}(u, d, a)$ be the reviewed set of material adverse rights-floor instances and define:

$$X_{\text{RF}}(u, d, a) = \sum_{\{k \in K_{\text{RF}}(u, d, a)\}} x_{\text{RF}}(k, r, a)$$

$$I_{\text{RF}}(u, d, a) = \tanh(\beta_{\text{RF}} \cdot X_{\text{RF}}(u, d, a))$$

The canonical default is $\beta_{RF} = \beta$, the governed saturation parameter used for direct welfare impacts in Section 5.5. In this rights-specific channel, I_{RF} and I_{rights} are evaluated separately for each protected right r ; the r subscript is suppressed in legacy display equations only. Machine records and validators MUST preserve the right identifier and MUST NOT merge overlapping rights into one compensable cell value. β_{RF} is a governed rights parameter, not an empirical constant. Every Tier 2-3 PCC MUST record it. Tier 3 and high-stakes Tier 2 runs MUST test the canonical value and at least one lower and one higher preregistered value whenever a plausible change could cross a rights threshold. Larger positive β_{RF} magnifies both positive and negative pre-saturation values toward the bounds and therefore makes a negative floor breach more likely for a fixed negative X_{RF} ; “tighter” MUST NOT be used without stating that direction. A non-default β_{RF} MUST be set ex ante, justified, applied symmetrically, and accompanied by the canonical-default counterfactual. If rights admissibility changes across the required sensitivity range, the run MUST report **RIGHTS_SATURATION_SENSITIVE** and MAY NOT claim an unqualified RF/NCRC result until governed review resolves the parameter basis. For the subgroup set $G_{RF}(u,d,a)$ required by Section 7.3, define $I_{rights}(u,d,a) = \min_{\{g \in G_{RF}(u,d,a)\}} I_{RF}(u,d,a|g)$; when no governed subgroup split is material, the declared base population is the sole member of G_{RF} . Apply the remaining Section 7.3 visibility and evidence rules without attenuation. Low confidence, short duration, or positive welfare elsewhere cannot attenuate this rights-floor channel.

7.4.2 Categorical-prohibition channel (Normative)

For each material rights instance, declare $z_{k,r} \in \{0,1\}$, where $z_{k,r} = 1$ means the option authorizes, requires, or knowingly imposes a categorical rights violation under the declared rights profile. Examples may include torture, intentional killing outside a lawful protective exception, non-consensual invasive medical intervention, arbitrary detention, deliberate deprivation below a basic-needs floor, or another prohibition explicitly declared by the governing profile.

A categorical flag is profile-governed and challengeable. It MUST NOT be inferred merely from rhetoric, and it MUST NOT be removed because the event is brief, affects a small protected subgroup, or produces benefits elsewhere.

7.4.3 Severe rights-hazard channel (Normative)

Explicit screening record (Normative). For every adverse rights-covered instance, the PCC MUST contain an explicit severe-rights-hazard screening field that records **ACTIVE**, reviewed **NOT_MATERIAL**, or **UNKNOWN**, together with the event-class and threshold basis, evidence source, reviewer status, and reopen trigger. An absent or blank screen is **UNKNOWN**, not **NOT_MATERIAL**. A run MUST NOT obtain RF/NCRC passage merely by omitting EventClass, h_{sev} , p^U , the exposed-subgroup denominator, or another field needed to determine whether the channel activates.

The severe-rights-hazard channel is distinct from the ordinary rights-floor threshold. It is activated by either a governed event class or a separate hazard-activation threshold. **Decision-materiality screen.** The background existence of a severe event class in a domain does not by itself activate the channel for every option. The run MUST determine whether the option creates, increases, transfers, conceals, controls, or makes a material claim about the protected population's exposure relative to the declared baseline. A reviewed **NOT_MATERIAL** finding requires an option-specific rationale, evidence source, reviewer status, and reopen trigger; it MUST NOT be represented as proof that background risk is zero.

For impact instance k and right r , define:

$SevereRightsHazardActive(k,r) := (EventClass_k \in C_{severe,r}) \text{ OR } (h_{sev,k,r} \geq \eta_r)$.

Canonical default: $\eta_r := \min(0.50, |\theta_r|)$. Domain profiles MAY tighten η_r but MUST NOT set it above 0.50 or disable a canonical event-class trigger.

$C_{severe,r}$ includes, where applicable: credible fatality pathways; torture or cruel, inhuman, or degrading treatment; serious bodily injury or forced invasive intervention; arbitrary detention or confinement; deprivation

of basic subsistence; collapse of due process or remedy; severe information coercion or manipulation; and irreversible or severe ecological destruction affecting protected life-support conditions.

When the channel is active, the run MUST separately declare intrinsic severity conditional on occurrence, likelihood or scenario probability, reversibility, consent/authority status, worst-affected subgroup, evidence status, and the governed tolerance or prohibition. The ordinary temporally weighted welfare product does not determine whether the hazard is reviewed. Missing gate-critical evidence produces NCRC_UNKNOWN, escalation, or refusal rather than an implied pass.

7.4.4 Violation components and predicate (Normative)

For each right r , define:

$$v_r^{\text{floor}}(a) = \max_{(u,d \text{ in } C_r)} (\theta_r - I_{\text{rights}}(u,d,a))^{+}$$

$$v_r^{\text{cat}}(a) = \max_{(k \text{ in } K_r(a))} z_{k,r}$$

$$v_r^{\text{risk}}(a) = \max_{(k \text{ in } K_r(a): \text{SevereRightsHazardActive}(k,r))} (p_{k,r}^{\wedge U} - \rho_{\{k,r\}})^{+}$$

where $(x)^{+} = \max(x, 0)$.

Reviewed-empty-set convention (Normative). For the categorical and severe-risk channels, $\max \emptyset = 0$ only after the run has completed the required stakeholder, event-class, and evidence review and has established that no material instance belongs in the applicable set. This convention does not apply when a required field, stakeholder, event class, or evidence trace is missing or unresolved; such cases produce NCRC_UNKNOWN, escalation, or refusal rather than zero.

If a required severe-hazard field is missing, set the rights evidence status to UNKNOWN rather than zero.

Rights-risk tolerance provenance rule (Normative). ρ_r is a governed maximum tolerated probability for an already activated severe-rights hazard; it is not inferred from the ordinary floor θ_r , analyst convenience, or the option's welfare score. Every non-zero ρ_r MUST identify the governing authority or domain profile, protected population and event class, evidence date, uncertainty basis, challenge/appeal path, and review trigger. Categorical prohibitions ordinarily use $\rho_r=0$. If ρ_r is missing, stale, unauthorized, or contested at a gate-material point, the severe-risk channel is NCRC_UNKNOWN, not zero. For each tested probability constraint, the tolerance MUST specify the protected bearer or population, adverse-event family, time horizon or exposure unit, and per-pathway or aggregate interpretation. Write the applicable tolerance as $\rho_{\{k,r\}}$; ρ_r in prose abbreviates that scoped value. For an aggregate tolerance, group all material pathways to the same protected event within that boundary and test a justified joint or union probability bound before taking the maximum violation. The risk-check index k includes these required aggregate probability records as well as individually constrained pathways; aggregate records are not additional welfare contributions. Independence, exclusivity, shared causes and repeated exposure require explicit treatment; a conservative union bound may be used where warranted. Do not assume independence or sum overlapping event probabilities as if disjoint. Keep individual constraints as additional checks. An unbounded material aggregate exposure or unresolved tolerance boundary yields NCRC_UNKNOWN, not passage.

Activation/computation examples (Normative). Failing case: for LIFE, let $\theta_{\text{LIFE}}=-0.90$, $\eta_{\text{LIFE}}=0.50$, intrinsic fatality-pathway severity $h_{\text{sev}}=0.80$, conservative probability bound $p^{\wedge U}=0.03$, and governed tolerance $\rho_{\text{LIFE}}=0.01$. The hazard activates because $0.80 \geq 0.50$, even though $0.80 < 0.90$; therefore $v_{\text{LIFE}}^{\text{risk}}=(0.03-0.01)^{+}=0.02$ and ordinary NCRC passage is blocked. Passing case: under an authorized domain profile for the same event class, let $p^{\wedge U}=0.004$ and $\rho_{\text{LIFE}}=0.005$, with complete provenance, conservative uncertainty, independent review, and a reopen trigger. Then $v_{\text{LIFE}}^{\text{risk}}=(0.004-0.005)^{+}=0$; the severe-risk channel passes, but the floor, categorical, evidence-completeness, TRC, CSV, and authority conditions remain independently binding. These examples do not establish a universal ρ_r ; tolerances remain right-, population-, event-, domain-, and authority-specific.

Three-channel worked illustration (Informative; equations normative in Section 7.4.4).

The three channels are non-substitutable: a pass on one channel cannot compensate a violation on another, and an unresolved material field yields NCRC_UNKNOWN rather than zero.

Right / channel	Declared input	Computed component	Interpretation
LIFE floor	$I_RF = -0.88$; $\theta_{LIFE} = -0.90$	$v_floor = \max(-0.90 - (-0.88), 0) = 0.00$	The floor channel passes.
LIFE categorical	No reviewed categorical prohibition instance	$v_cat = 0.00$	Zero follows only after the required reviewed-empty-set check.
LIFE severe-risk	$p^U = 0.03$; $\rho_{LIFE} = 0.01$; severe hazard active	$v_risk = 0.02$	NCRC fails despite the floor-channel pass.
DIGN floor	$I_RF = -0.70$; $\theta_{DIGN} = -0.55$	$v_floor = 0.15$	A direct floor breach independently fails NCRC.

NCRC(a) = TRUE if and only if, for every right r:

- $v_r^{floor}(a) = 0$;
- $v_r^{cat}(a) = 0$;
- $v_r^{risk}(a) = 0$; and
- no material rights evidence field is unresolved.

The NCRC-passing set is:

$$A_NCRC = \{a \text{ in } O : NCRC(a) = \text{TRUE}\}$$

Interpretation: a short-duration or low-confidence rights violation cannot pass merely because welfare-style temporal or confidence multipliers are small. Rights-floor impact, categorical prohibition, and severe-hazard exposure are tested separately and non-compensatorily.

Rights-conflict precondition (Normative). Before Rights Emergency Mode is invoked, the run MUST attempt and document feasible alternative generation, redesign, scope narrowing, sequencing, delay, mitigation, or other means of producing at least one RF/NCRC-passing option. An unresolved rights conflict remains a structured rights-vector problem; it MUST NOT be converted into an ordinary welfare trade-off. Only when no rights-passing option remains may Section 7.5 be used, subject to necessity, least-rights-infringing means, time limitation, independent challenge, remedy, and return-to-normal requirements.

7.4.5 RightsEffectToken Non-Dilution (Normative)

Every decision-material rights harm SHALL carry one RightsEffectTokenID binding the affected bearer or exposed subgroup, protected right, state, option, floor baseline, causal pathway, maximal contiguous time window, and evidence object.

RF/NCRC SHALL evaluate the full unallocated token once before any analytical Union-Scope allocation. Allocation, deduplication, scope tagging, dimension routing, or emergent-scale analysis MAY affect residual RLS accounting but MUST NOT reduce the token-level rights calculation.

For rights token j, right r, and option a, define $X_RF^{token}(j,r,a) = \mu_j \cdot r \cdot j^R \cdot \ell_j^U \cdot e_j$, with $\tau_RF=1$, $c_gate=1$ for adverse gate computation, and $h_base=1$. Define $I_RF^{token} = \tanh(\beta_RF \cdot X_RF^{token})$ and $v_{j,r}^{token} = \max(0, \theta_r - I_RF^{token})$ under the sign convention and threshold orientation of Appendix C.

The controlling floor-channel violation is the maximum of the existing cell-level violation and every applicable unallocated token-level violation. Categorical and severe-hazard channels remain independently required.

If one token is presented through multiple sub-threshold scope views, or token identity is missing where allocation could affect the result, set RIGHTS_PARTITION_SENSITIVE or RIGHTS_EFFECT_TOKEN_MISSING. The option MUST NOT receive RF_PASS until the unallocated token has been evaluated and all required fields are resolved.

Distinct harms to different bearers remain distinct tokens. This rule prevents dilution of one harm; it does not sum unrelated small harms into a synthetic super-harm.

Rights non-dilution record	Required content
RightsEffectTokenID	Stable identifier across every analytical view
Bearer / exposed subgroup	Who bears the protected harm; reach is evaluated at exposed-subgroup level
Protected right and coverage route	Applicable right, primary covered cell, extension/routing if needed
Unallocated token calculation	Full harm before scope allocation
Secondary scope views	Visible but non-attenuating
Disposition	PASS, categorical/risk/floor violation, UNKNOWN, or required redesign/escalation

7.5 Emergency Mode (Rights-Failure Handling) (Normative)

If $A_NCRC = \emptyset$ (no option passes rights), RippleLogic enters Emergency Mode. Emergency Mode is not a loophole; it is a controlled failure protocol.

7.5.1 Emergency Mode rights priority order (Normative)

Rights are ordered lexicographically by the canonical emergency default priority:

[LIFE, BODY, ECOL, LBTY, NEED, DIGN, PROC, INFO]

This order is a provisional governance default for constrained emergency comparison, not a discovered universal moral ranking. It prioritizes immediate survival and bodily integrity, then ecological conditions capable of undermining life-support systems, while retaining liberty, need, dignity, procedure, and information as independently non-compensatory rights outside Emergency Mode. The default MUST be used for the canonical result, but every Tier 3 or high-stakes emergency run MUST also test at least one defensible alternative ordering that moves PROC and INFO earlier when procedural or epistemic failure could causally enable LIFE, BODY, LBTY, or NEED harms. If the selected option changes under a defensible alternative order, RippleLogic MUST report EMERGENCY_RIGHTS_ORDER_SENSITIVE, refuse to portray the result as uniquely entailed by the framework, and route the decision to lawful accountable authority with independent challenge and the complete rights vector disclosed.

7.5.2 Emergency Mode selection rule (Normative)

Construct a compound violation tuple for each right r :

$$q_r(a) = (v_r^{\text{cat}}(a), v_r^{\text{risk}}(a), v_r^{\text{floor}}(a))$$

and the option vector:

$$q(a) = (q_LIFE(a), q_BODY(a), q_ECOL(a), q_LBTY(a), q_NEED(a), q_DIGN(a), q_PROC(a), q_INFO(a))$$

Select the option that lexicographically minimizes $q(a)$, so categorical prohibition is considered before severe-hazard excess and floor depth within each right. If tied:

- Minimize TRC CVaR (if TRC scenarios exist or are rapidly constructed), then
- If at least one tied emergency candidate passes TRC, compare residual welfare only among those tied TRC-passing candidates. If every tied candidate also fails TRC, apply Section 8.8 Tail Emergency Mode and do not use ordinary RLS ranking.

7.5.3 Independent Challenger Requirement (Normative)

Before Emergency Mode can be invoked, an independent challenger MUST propose at least one alternative option.

Independence requirements (MUST):

- No reporting relationship to decision owner
- No material interest in decision outcome
- Access to the same information set

Minimum challenger effort:

- Tier 2: ≥ 30 minutes active option generation
- Tier 3: ≥ 2 hours active option generation

If time pressure prevents this, PCC MUST record CHALLENGE_DEFERRED_EMERGENCY and a retrospective challenge review MUST occur within 24 hours (or earliest feasible time) with an addendum PCC.

7.5.4 Emergency documentation requirements (Normative)

PCC MUST include:

- Emergency declaration and trigger conditions
- $v_r(a)$ for each option
- Lexicographic comparison trace
- Mitigation/remediation plan
- Review cadence and return-to-normal triggers

7.6 Anti-Gaming Rules Specific to NCRC (Normative)

- Rights-covered cells MUST NOT be masked out of RLS aggregation if they are non-maskable by policy (Appendix H audit flags).
- Rights checks MUST use worst-off subgroup semantics and must not be evaluated on averages.
- “Unknown” data MUST NOT be treated as neutral; missing data triggers ignorance penalties (Section 5.7).
- Emergency Mode MUST NOT be used without challenger protocol and remediation plan.

SECTION 8: CATASTROPHIC RISK LAYER: TAIL-RISK CONSTRAINT (TRC)

8.0 Purpose and UBE Rationale (Normative)

Even when rights floors are respected in expectation, a decision can carry non-trivial probability of catastrophic harm to Humanity / Global Coordination or the Biosphere. Under UBE, catastrophe avoidance is lexicographically prior to welfare optimization because catastrophic outcomes can eliminate future choice across unions and collapse enabling conditions.

TRC therefore bounds catastrophic exposure using CVaR, a coherent tail-risk measure (Artzner et al., 1999; Rockafellar & Uryasev, 2000). This posture is motivated by ruin dynamics under deep uncertainty (Taleb, 2012) and by Earth-system research indicating that destabilization of life-support conditions can occur through tipping cascades and boundary transgression corridors with multi-scale downstream harms (Rockström et al., 2009, 2023; Steffen et al., 2015; IPCC, 2023).

Modeled-safety interpretation (Normative clarification). Passing TRC demonstrates compliance with the declared scenario library, probability policy, catastrophe-cell construction, and corridor threshold for the run. It does not prove empirical invulnerability outside the modeled scenario set. Implementations **MUST NOT** present TRC passage as a guarantee that no catastrophic pathway exists; it is a governance-grade bound within the declared scenario model.

8.1 Tier Requirements for TRC (Normative; strict posture)

Tier	TRC Required?	Minimum
Tier 1	Required qualitative trigger screen	N/A
Tier 2	REQUIRED when catastrophe relevance plausible	≥5 minimum
Tier 3	REQUIRED	≥20 minimum
Tier 4	Design target only (claim prohibited in this release line)	N/A

8.1.1 Catastrophe relevance trigger (Normative)

Catastrophe relevance is plausible if any of the following are true:

- The decision materially affects any catastrophe cell in C_cat (Section 8.2), OR
- The decision has plausible pathways to tipping cascades (climate, bio, conflict, infrastructure, financial), OR
- The decision is materially irreversible or creates lock-in at Polity/CMIU/Biosphere scales, OR
- Credible challengers identify a plausible catastrophic failure mode even if decision owners prefer not to model it.

If triggered, TRC **MUST** be executed at the tier's required rigor.

TRC not-triggered state (Normative). If a documented catastrophe-relevance assessment addresses every Section 8.1.1 trigger and finds none plausible within the declared decision and evidence boundary, the run records TRC_NOT_TRIGGERED. This is a pass-equivalent qualification state for cascade routing only. It **MUST NOT** be described as TRC_PASS, as a CVaR result, or as evidence that catastrophe is impossible. The assessment, evidence boundary, reviewer status, and reopen triggers **MUST** be recorded. New catastrophic, irreversible, lock-in, or ruin-relevant evidence reopens TRC before RLS.

8.2 Catastrophe Cell Set C_cat (Normative)

Base catastrophe cell set:

$C_cat_base = \{(6, \text{Health}), (6, \text{Environment}), (7, \text{Environment})\}$

These correspond to:

- Humanity / Global Coordination-Health: Global-scale health viability (pandemic mortality, mass disability, collapse of health capacity)
- Humanity / Global Coordination-Environment: Environment-as-civilization-condition (habitability, agricultural stability, climate-driven displacement)
- Biosphere-Environment: Earth-system integrity (planetary boundaries, biodiversity, biogeochemical stability)

The base set explicitly includes biosphere integrity and civilization-scale habitability conditions because “safe operating space” and “safe and just boundaries” analyses identify global environmental corridors as central determinants of long-run human and ecological viability (Rockström et al., 2009, 2023).

Extensions: Tier 1-3 **MAY** extend C_cat only with explicit PCC justification and **MUST** apply the same set to all compared options.

Non-dilution rule for catastrophe profiles (Normative). Adding an extension **MUST NOT** lower the measured loss or improve the TRC disposition for an unchanged base catastrophe merely because normalized weight is spread across more cells. The base catastrophe profile and every triggered extension

profile SHALL therefore be normalized and evaluated separately. Compute $L_p(a,s)$ and $CVaR_{\alpha,p}(a)$ for each active profile p ; ordinary TRC qualification requires every active profile to satisfy its governed corridor. A combined summary MAY be reported, but it cannot rescue a failed profile. Any alternative non-dilutive construction requires ex ante governance, proof that unchanged base losses cannot decrease when an extension is added, symmetric application across options, and PCC disclosure.

Clarification (Normative). C_{cat_base} is the default habitability/civilization catastrophe proxy. It cleanly captures biosphere degradation and civilization-scale health/environment collapse. Governance lock-in, coercive surveillance, procedural collapse, information-ecosystem capture, or other agency-destroying catastrophes are catastrophe-relevant and SHALL NOT be treated as rhetorically covered by the base triplet alone. When any governance-lock-in trigger is plausible, including surveillance normalization, emergency-power persistence, due-process suspension, information-ecosystem capture, irreversible dependency lock-in, or concentrated arbitrary control over critical infrastructure, the run MUST activate the default governance-lock-in extension profile defined in Appendix D.1A, unless a stricter governed profile is declared in the PCC. The chosen profile MUST be applied symmetrically across all compared options and recorded in the PCC.

WTSL-AIX / governance-lock-in tail category is mandatory when AI deployment or governance-lock-in triggers are plausible, including surveillance normalization, emergency-power persistence, due-process suspension, information-ecosystem capture, irreversible dependency lock-in, or concentrated arbitrary control. $AI_DEPLOYMENT=YES$ is sufficient but not necessary.

8.3 Scenario Governance (Normative)

Scenario-discovery record (Normative). Tier 3 and high-stakes or decision-material Tier 2 runs MUST identify the scenario-discovery method or combination of methods, material sources, participating experts or affected stakeholders, independent challenger or red-team contribution where required, search boundary, and disposition of materially plausible rejected scenarios. Suitable methods may include structured elicitation, domain hazard analysis, incident and near-miss review, historical analogues, affected-stakeholder input, literature and standards review, and AI-assisted candidate generation. Generated or institutionally supplied scenarios are discovery inputs; they are not self-validating probability estimates or proof that the scenario surface is complete.

TRC is evaluated over a governed scenario set S with probabilities p_s such that $p_s \geq 0$ and $\sum_s p_s = 1$.

8.3.1 Mandatory Tail Scenario Categories (Normative)

The scenario set MUST include the following tail categories unless explicitly justified as not plausible for the decision context. Governance / lock-in catastrophe is a mandatory tail category if and only if the governance-lock-in extension profile in D.1A is active; otherwise it is recommended but not mandatory. When that profile is active, the scenario set MUST include at least one governance-lock-in tail category with category probability floor $p_{floor} \geq 0.02$.

- Pandemic/biological disruption
- Climate tipping cascade
- Financial system collapse
- Major conflict escalation
- Critical infrastructure failure

Governance / lock-in catastrophe

Minimum probability floor per category: $p_{floor} \geq 0.02$

This prevents “include but set near-zero probability” gaming.

Probability-floor epistemic boundary. The p_{floor} value is an anti-omission governance prior, not a calibrated empirical frequency estimate. It prevents mandatory tail classes from being made invisible by arbitrary near-zero assignment and must be sensitivity-tested where decision-material.

Scenario partition and dependence rule (Normative). Tail categories are coverage labels, not automatically mutually exclusive events. The PCC MUST declare whether scenarios form a mutually exclusive partition, overlap, or belong to dependence clusters. Overlapping category floors MUST NOT be summed as though independent. Tier 3 and decision-material Tier 2 runs MUST either construct mutually exclusive scenarios, use bounded probability intervals/sets, or run a robust CVaR sensitivity over plausible probability assignments. Correlated compound scenarios and shared causal drivers MUST be tagged so probability mass is not double-counted and independence is not silently assumed.

Rejected-scenario log. Tier 3 runs and high-stakes Tier 2 runs MUST record any material tail category or challenger scenario rejected as implausible, including the rejection rationale, evidence basis, reviewer status, and whether the rejection would change if the reality surface, probability floor, or time horizon changed.

These mandatory tail categories are not intended as exhaustive; they are a minimal governance floor reflecting cross-domain systemic risk exposure in tightly coupled global systems, including climate-driven boundary cascades (Rockström et al., 2009, 2023; Steffen et al., 2015; IPCC, 2023).

8.3.2 Implausibility test for exempting mandatory tail categories (Normative)

If a deployment wishes to exempt a mandatory tail category as “implausible,” the exemption requires:

- (i) Written justification in the PCC explaining why the category is implausible for this specific decision context
- (ii) Independent reviewer sign-off (Tier 3 required; Tier 2 recommended)
- (iii) Recording in PCC with audit flag MANDATORY_TAIL_CATEGORY_MISSING_WITH_JUSTIFICATION

The burden of proof is on omission, not inclusion. This preserves UBE.

8.3.3 Scenario set size minima (Normative)

Tier	Minimum
Tier 2 (when TRC required)	≥5 minimum
Tier 3	≥20 minimum

If a tier cannot meet its minimum scenario count, the run MUST NOT claim tier compliance for TRC coverage and MUST record audit_flag SCENARIO_LIBRARY_MIN_EXCEPTION in the PCC, including written justification and a remediation plan to reach the minimum scenario count for the next comparable run.

8.4 Loss Construction (Tier 1-3: Bounded-Impact TRC) (Normative)

Tier 1-3 TRC uses bounded-impact loss based on propagated Base-stream impacts in catastrophe cells.

Domain constraint: All propagated impacts are bounded: $I_{prop_base,c}(a|s) \in [-1, +1]$.

For each scenario s , define scenario loss:

$$L(a,s) = \sum_{c \in C_cat} \omega_c \times \max(0, -I_{prop_base,c}(a|s))$$

Normative stream binding (Design A): TRC MUST be computed from the Base stream (I_{prop_base}). TRC MUST NOT use Welfare-stream (sentience-weighted) impacts.

Where:

- $\omega_c \geq 0$, $\sum_{c \in C_cat} \omega_c = 1$ (catastrophe weights are normalized)
- $\max(0, x)$ returns x if $x > 0$, otherwise 0 (only negative impacts contribute)
- $I_{prop_base,c}(a|s)$ is the propagated impact in catastrophe cell c under scenario s

Default catastrophe weights (if not otherwise governed): Uniform weights over C_cat .

Catastrophe-weight governance (Normative). Any non-default ω_c vector MUST be declared ex ante in the PCC, normalized within the active catastrophe profile, justified by a named evidence or governance source,

and applied symmetrically across compared options and scenarios. It MUST NOT weaken the base-profile non-dilution rule, hide a material catastrophe cell, or convert an unchanged severe loss into a pass merely through weight spreading. Tier 3 and decision-material or high-stakes Tier 2 runs using non-default ω_c MUST report a uniform-weight counterfactual and any material sensitivity in loss, CVaR, severity-class disposition, or TRC status.

TRC loss-reference rule (Normative). In this edition, $I_{\text{prop_base},c}(a|s)$ represents the scenario-conditioned residual condition under option a relative to the declared safe-corridor reference for catastrophe cell c , after supported controls and the applicable GateAdverseBound. It is not merely the change from a potentially unsafe no-action outcome. Use the same declared exposure domain, reference, horizon and catastrophe-weight profile across compared options. The domain must include material affected stakeholders, dependencies and foreseeable exposure; any exclusion needs evidence and challenger-visible rationale and cannot be chosen to manufacture a pass. Exogenous scenarios remain required where applicable, but their inclusion does not assert that every unrelated planetary condition belongs to every local decision. Report same-scenario continuation and option-attributable change separately when they aid causal interpretation; subtracting continuation loss MUST NOT erase residual unsafe exposure. No action, delay and continuation receive the same qualification tests. A pre-existing in-scope corridor breach can make every ordinary option fail; then use the existing redesign, escalation or explicitly qualified emergency procedures, never an automatic zero-loss baseline. Option-varying scenario losses are compatible with this residual-state interpretation and do not, by themselves, establish incremental-only semantics.

Interpretation:

- Only negative impacts in catastrophe cells contribute to loss.
- Loss is bounded in $[0,1]$ because weights are normalized and impacts lie in $[-1,+1]$.
- Positive impacts in catastrophe cells (improvements) do not offset losses; they simply contribute zero to the loss function.

Example: If $C_{\text{cat}} = \{(6, \text{Health}), (6, \text{Environment}), (7, \text{Environment})\}$ with uniform weights ($\omega_c = 1/3$ each), and under scenario s the propagated impacts are:

- $I_{\text{prop_base}}(6, \text{Health}, a|s) = -0.30 \rightarrow \text{contributes } 1/3 \times 0.30 = 0.10$
- $I_{\text{prop_base}}(6, \text{Environment}, a|s) = +0.10 \rightarrow \text{contributes } 1/3 \times \max(0, -0.10) = 0$
- $I_{\text{prop_base}}(7, \text{Environment}, a|s) = -0.60 \rightarrow \text{contributes } 1/3 \times 0.60 = 0.20$

Then $L(a,s) = 0.10 + 0 + 0.20 = 0.30$

8.4A Catastrophe Severity-Class Overlay (Normative)

Bounded loss supports stable computation but can compress materially different catastrophe classes near the same ceiling. Every TRC scenario MUST therefore declare one severity class in addition to $L(a,s)$:

C1_SEVERE, **C2_SYSTEMIC**, **C3_CIVILIZATION_THREATENING**, or **C4_EXISTENTIAL**. The class is evidence-governed and challengeable; it is not inferred from rhetoric or from the bounded loss number alone.

- **C1_SEVERE**: large but bounded harm with functioning recovery institutions.
- **C2_SYSTEMIC**: multi-system or cross-polity disruption with major recovery burden.
- **C3_CIVILIZATION_THREATENING**: plausible loss of durable civilizational capacity, global coordination, or planetary habitability conditions.
- **C4_EXISTENTIAL**: plausible permanent destruction of humanity or comparable foreclosure of future moral patients.

A governed profile MAY set class-specific probability caps or stricter corridors. Until such a profile is declared, any material **C4_EXISTENTIAL** pathway prevents an ordinary TRC pass unless its conservative upper probability bound is shown to be below an explicitly governed existential-risk tolerance. A

C3_CIVILIZATION_THREATENING pathway requires independent challenge, class-specific monitoring/controls, and sensitivity to a stricter corridor. This is an overlay inside TRC, not a sixth gate, and it

prevents a bounded scalar from erasing the qualitative difference between severe, systemic, civilization-threatening, and existential outcomes.

8.5 CVaR Computation (Normative)

Let $L(a)$ be the random loss induced by scenarios s with probabilities p_s .

Value-at-Risk: $\text{VaR}_\alpha[L(a)] = \inf\{z : P(L(a) \leq z) \geq \alpha\}$

CVaR operational definition (Tier 1-3). For RippleLogic Tier 1-3 discrete scenario libraries, $\text{CVaR}_\alpha[L(a)]$ is the probability-weighted average loss over the worst $\beta = 1 - \alpha$ probability mass, computed by the canonical discrete algorithm in Appendix D.7.

Continuous-case shorthand. In continuous loss distributions without probability mass at VaR, this corresponds to the familiar expression $E[L(a) \mid L(a) \geq \text{VaR}_\alpha[L(a)]]$. In discrete scenario sets with atoms at the VaR boundary, that shorthand is not controlling and **MUST NOT** override Appendix D.7.

The Appendix D.7 discrete tail-mass algorithm is normative, audit-replayable, and controlling for all finite Tier 1-3 scenario sets, including partial mass at the boundary scenario when required.

8.6 TRC Admissibility Predicate (Normative)

An option passes TRC if and only if:

$\text{CVaR}_\alpha[L(a)] \leq \tau_{\text{TRC}}$, and every active class-specific cap or tolerance under Section 8.4A is satisfied, with no unresolved material C3 or C4 field.

Define:

$A_{\text{adm}} = \{a \in A_{\text{NCRC}} : \text{TRC}(a) = \text{TRUE}\}$

8.7 Default TRC Parameters by Context (Normative Defaults)

These are defaults; governance may tighten them. Loosening requires explicit governance justification and PCC recording.

Context	α (tail level)	τ_{TRC} (corridor threshold)
Personal	0.90	0.30
Organizational	0.95	0.20
Reversible policy	0.95	0.15
Irreversible policy	0.99	0.10
Existential risk	0.999	0.05

High-alpha effective-maximum disclosure (Normative clarification). With discrete scenarios, CVaR_α equals the worst modeled loss whenever the worst-loss scenario probability is at least the tail mass $1 - \alpha$. Under the mandatory scenario probability floor, this will commonly occur for $\alpha = 0.99$ or 0.999 . This is an intended conservative limit, not an averaging result. The PCC **MUST** report the effective tail support, whether CVaR is numerically identical to the modeled maximum loss, and the consequence of that equivalence for the TRC result. A maximum-equivalent CVaR remains bounded by the declared scenario library and does not establish that every catastrophic pathway has been identified.

Low-alpha corridor interpretation (Informative clarification). In lower-stakes contexts such as the Personal default, the scalar CVaR corridor can be permissive for a very-low-probability severe modeled loss at the mandatory probability floor. The Section 8.4A severity-class overlay remains a separate conjunctive backstop:

a material C3/C4 pathway does not become an ordinary pass merely because the scalar corridor passes. Where the overlay, rather than the scalar corridor, controls disposition, the PCC SHOULD say so explicitly. This clarification changes no α , τ_{TRC} , p_{floor} , or severity-class rule.

8.7A TRC Trigger Challenge and High- α Limit (Normative Clarification)

High-stakes Tier 2 and Tier 3 uses of TRC_NOT_TRIGGERED require an independent challenger disposition or documented infeasibility of independent challenge. Without that disposition, the strongest permitted claim is ASSUMPTION_BOUND and the run SHALL preserve a reopen trigger.

Discrete CVaR procedure. Normalize scenario probabilities, sort losses from highest to lowest, identify the boundary scenario at which cumulative probability reaches tail mass $1-\alpha$, take all probability mass above the boundary, take only the exact fraction of the boundary scenario required to fill the tail, and divide the resulting loss mass by $1-\alpha$. The sorted table, boundary index, partial boundary mass, and final value SHALL be retained.

High- α limit. When $1-\alpha$ is smaller than the probability assigned to the worst modeled scenario, $CVaR_\alpha$ equals the worst modeled loss. At that setting TRC acts as a maximum-modeled-loss screen; probability improvements inside that worst scenario do not change CVaR until the tail boundary changes. This is conservative and SHALL be disclosed.

This clarification does not recalibrate α , p_{floor} , τ_{TRC} , severity classes, or corridor values. Any recalibration requires a versioned reference suite, sensitivity analysis, and independent review.

8.8 Tail Emergency Mode (Normative; replaces automatic TRC fallback)

If A_{NCRC} is non-empty but no rights-admissible option passes TRC, the ordinary selectable set is empty. The default response is redesign, delay, escalation, no-action review, or refusal. A least-CVaR option MUST NOT be automatically promoted into ordinary selection.

Tail Emergency Mode MAY be invoked only when the PCC demonstrates all of the following:

1. redesign and materially safer alternatives have been actively attempted;
2. delay, refusal, or no-action is unavailable or is evidenced to create equal or greater ruin exposure;
3. an independent challenger has reviewed the emergency claim and proposed alternatives;
4. a governed absolute maximum exposure τ_{TRC_max} has been declared and the provisional option does not exceed it;
5. the action is time-bounded, scope-bounded, reversible where possible, and paired with mitigation, monitoring, shutoff, remedy/redress, and explicit exit criteria;
6. one-tier-higher approval and lawful authority are recorded; and
7. the decision will be reopened on a declared review cadence or immediately when a safer option becomes available.

Within the emergency candidate set that satisfies these conditions, choose the provisional option with minimum $CVaR_\alpha$. Record $TAIL_EMERGENCY_MODE_INVOKED$, the CVaR deficit above τ_{TRC} , the absolute-cap margin, the necessity evidence, challenger record, controls, burden bearers, and return-to-normal trigger.

The result is $TAIL_EMERGENCY_PROVISIONAL_ACTION$. It is not TRC_PASS, SELECT, aligned, safe, or ordinarily admissible. Ordinary RLS ranking is prohibited in Tail Emergency Mode.

If no candidate remains below τ_{TRC_max} , RippleLogic MUST refuse the stronger action claim unless a separately authorized rights-and-tail emergency protocol with stricter public-law or domain controls governs the case. If both NCRC and TRC fail for every option, Section 7.5 rights priority governs first, while the Tail Emergency conditions in this section remain binding for any option that also exceeds the ruin corridor.

8.9 Anti-Gaming Requirements for TRC (Normative)

- Mandatory tail categories cannot be omitted without explicit implausibility justification per Section 8.3.2.

- Probability floors apply per category; they cannot be silently violated.
- Scenario sets must be comparable across options (same S, same p_s).
- PCC must record scenario provenance and how probabilities were assigned.
- Tier 2-3: if probabilities are highly uncertain, the PCC SHOULD include sensitivity over p_s (Tier 3 required).

SECTION 9: CSV GATE: CONTAINMENT AND STRUCTURAL VIABILITY

9.0 Purpose and UBE Rationale (Normative)

CSV means Containment and Structural Viability. CSV is the strengthened fourth public level in the five-level RippleLogic cascade inside the MathGov framework:

RG -> RF -> TRC -> CSV -> RLS

Public-language mnemonic:

Ground -> Protect -> Bound -> Contain/Verify -> Score

CSV prevents a core failure mode under Union-Based Reality: a local option can look beneficial for a sub-union while quietly degrading the larger system, capacity, relation, ecological support, institutional trust, or future choice-space that makes the benefit possible. Under Union-Based Ethics, such “gains” are not counted as system-level improvement when they materially degrade, hollow out, overload, or exploit their containing systems beyond tolerance.

CSV is therefore the selectability gate applied after NCRC and TRC and before final RLS ranking. CSV is not a sixth public level. It is the upgraded fourth level, replacing the older public label “Containment” with the more complete label “Containment and Structural Viability.” Structural Viability is no longer merely a quiet subcheck; it is explicitly visible inside CSV.

Core question:

Can this option deliver its claimed good without secretly breaking, hollowing, overloading, exploiting, or depending on the systems that make the good possible?

Plain-language rule:

CSV is not a purity filter. CSV is a reality-integrity filter.

9.1 CSV Non-Purity Rule (Normative)

A negative ripple is not automatically a CSV failure. Every real action carries some cost, burden, risk, or tradeoff. CSV MUST distinguish between detected harm, material residual harm, conditionable harm, redesign-trigger harm, and gate-failing harm.

CSV SHALL NOT demand zero harm. CSV SHALL demand that harms are not hidden, unbounded, structurally degrading, rights-adjacent without review, irreversible without authorization, dependency-producing without transition, or unjustly externalized onto affected unions.

Status categories:

Status	Meaning	Result
DETECTED_HARM	A negative effect exists.	Must be recorded and routed to the correct layer.
RESIDUAL_HARM	Harm is material but bounded, visible, and below CSV failure threshold.	Carries into RLS as a welfare cost.
CONDITIONABLE_HARM	Harm is acceptable only with controls.	CSV_PASS_WITH_CONTROLS if controls are

Status	Meaning	Result
		binding.
REDESIGN_TRIGGER	Harm is too high as specified but plausibly fixable.	CSV_REDESIGN_REQUIRED; option is not selectable as specified.
CSV_FAILURE	Harm is uncontained, structurally degrading, non-viable, unjustly externalized, or lock-in producing beyond tolerance.	CSV_FAIL; option is not selectable.

Normative consequence:

RLS may rank ordinary residual tradeoffs among selectable options. RLS MUST NOT rescue an option that fails CSV.

9.2 CSV Status Ladder (Normative)

CSV outputs one of the following status values:

CSV status	Meaning	Selection consequence
CSV_PASS	Harms are contained, structural viability is adequate, and residual burdens may be carried into RLS.	Option may proceed to RLS if NCRC and TRC passed.
CSV_PASS_WITH_CONTROLS	Option is selectable only if declared controls, mitigations, monitoring, or limits are binding.	Option may proceed to RLS as the controlled version only.
CSV_REDESIGN_REQUIRED	Option is not selectable as specified, but a redesigned version may be evaluated.	Original option excluded; redesigned option may re-enter RG/cascade.
CSV_ESCALATE	Stakes, uncertainty, missing evidence, or governance sensitivity require higher review.	No final selection until escalation is resolved.
CSV_FAIL	Option unacceptably damages, hollows, overloads, exploits, or depends on containing systems beyond tolerance, or is structurally non-viable.	Option is not selectable.
CSV_EMERGENCY_PROVISIONAL	Temporary selection under necessity where no better feasible option currently exists and strict emergency controls are declared.	Time-bounded provisional selection only; no unqualified alignment claim.
CSV_NOT_MATERIAL	CSV burden is not material for this run, with rationale.	Option may proceed if other gates pass, but rationale is recorded.

A CSV_PASS_WITH_CONTROLS is a pass only for the controlled option. If controls are removed, expire, fail, or are materially altered, CSV MUST be rerun.

A CSV_EMERGENCY_PROVISIONAL disposition MUST include a time limit, harm cap, monitoring plan, review trigger, and transition or remediation plan. It MUST NOT be marketed as full alignment, ordinary selectability, deployment certification, or Tier 4 validation.

9.2A CSV status resolver guide (Normative for Tier 3; recommended for Tier 2)

The following resolver constrains how evidence and diagnostics map into CSV statuses. It does not eliminate reviewer judgment, but reviewer judgment must be recorded against these routing rules.

Evidence / diagnostic condition	Required CSV routing
Any material CSV diagnostic is level 4, or the option creates uncontained, structurally degrading, unjustly externalized, non-viable, hidden, lock-in-producing, unmonitored, or beyond-tolerance harm	CSV_FAIL

Evidence / diagnostic condition	Required CSV routing
Any Structural Viability subcheck materially fails and the failure is not time-bounded, controlled, or redesignable	CSV_FAIL
Any Structural Viability subcheck materially fails but a plausible redesigned option could repair it	CSV_REDESIGN_REQUIRED
Gate-material CSV evidence is missing, stale, contested, or insufficient to support selectability	CSV_ESCALATE, claim-narrowing, or refusal of the stronger selection claim
Any material diagnostic is level 3	CSV_REDESIGN_REQUIRED or CSV_ESCALATE, with rationale
A material diagnostic is level 2 and binding controls, owners, monitoring, review triggers, and failure consequences are declared	CSV_PASS_WITH_CONTROLS
A material diagnostic is level 2 without binding controls	CSV_REDESIGN_REQUIRED
All material diagnostics are level 0-1 and structural viability subchecks pass	CSV_PASS
CSV burden is genuinely not material to the declared claim boundary and the PCC records why	CSV_NOT_MATERIAL
Necessity is documented, no better feasible option is available, controls are time-bounded, harm caps are declared, monitoring is active, and a transition/remediation path exists	CSV_EMERGENCY_PROVISIONAL

Resolver precedence (Normative clarification). Ordinary PASS, PASS_WITH_CONTROLS and NOT_MATERIAL rows cannot override an applicable FAIL, REDESIGN_REQUIRED, ESCALATE or unresolved-evidence route. Established non-viable or beyond-tolerance harm controls over permissive rows. Where redesign and evidence escalation both apply, record both requirements and the controlling disposition; neither permits ordinary selectability until resolved. CSV_EMERGENCY_PROVISIONAL is a separate prerequisite-bound exception and never a less restrictive ordinary pass.

CSV maturity boundary. CSV may be binding in Tier 3 run-level conformance using declared provisional structural indicators, but this does not imply cross-domain validated UCI/HOI measurement, ProofPack readiness, or machine-verifiable structural truth. UCI/HOI remain diagnostics unless and until separate validation artifacts establish stronger measurement status.

9.3 CSV Core Tools (Normative)

CSV uses five core tools. Implementations MAY add domain-specific tools, but MUST preserve the five core functions when material.

9.3.1 Containing-System Map

Identify the systems that hold, sustain, receive, or absorb the option's consequences. Relevant systems MAY include persons, households, communities, organizations, institutions, infrastructure, ecology, attention, trust, labor capacity, knowledge systems, public legitimacy, supply chains, financial reserves, legal authority, and future option-space.

Minimum questions:

- What local good is the option claiming?
- Which union receives the local gain?
- Which larger unions contain or sustain that gain?
- What capacities, relations, resources, and ecological supports does the option depend on?
- What systems receive spillover costs?
- Which affected groups can contest, consent, or recover?

9.3.2 Structural Viability Check

Structural Viability asks whether the option can actually stand in the declared reality surface. Ethical attractiveness is not enough. A desirable option may be a design target while still failing current execution viability.

Minimum subchecks:

Code	Subcheck	Question
SV-1	Resource closure	Are budget, time, people, compute, energy, materials, attention, and authority sufficient?
SV-2	Dependency closure	Are prerequisites satisfied or assigned to credible owners and timelines?
SV-3	Reversibility / endpoint coherence	Is rollback, endpoint, or irreversibility status declared and coherent?
SV-4	Capacity non-exceedance	Does execution stay within operational, institutional, technical, and enforcement capacity?
SV-5	Internal non-contradiction	Does the option avoid mutually incompatible requirements or declare a grounded reconciliation mechanism?

An option with material SVC failure is not selectable for current execution, regardless of RLS rank.

9.3.3 Externality Containment Ledger

CSV MUST make spillovers visible. For material runs, the PCC SHOULD include an ExternalityContainmentLedger with at least:

Field	Required content
Spillover	Smoke, noise, labor burden, dependency, extraction, attention harm, ecological stress, etc.
Receiver	Which union, subgroup, system, or future condition receives the burden.
Severity	Normalized severity or qualitative tier.
Duration	Acute, temporary, recurring, chronic, permanent, or irreversible.
Mitigation	Controls, design changes, monitoring, compensation, substitution, or remediation.
Residual	What burden remains after mitigation.
Routing	NCRC, TRC, CSV, RLS, monitoring, or escalation.

Unrecorded material externalities are not neutral. They create CSV uncertainty and may trigger CSV_ESCALATE or CSV_REDESIGN_REQUIRED.

9.3.4 Substitution Ladder

CSV MUST ask whether a less harmful feasible alternative exists. The ladder is not a requirement to choose the ideal option immediately; it is a discipline against convenience-based harm.

Default substitution ladder:

1. Avoid the harmful need.
2. Reduce the demand.
3. Use low-harm alternative.
4. Use hybrid/backup alternative.
5. Use higher-harm option only with controls.
6. Use emergency provisional option only under necessity and review.
7. Reject or redesign if harms cannot be contained.

If a materially cleaner and feasible alternative exists, choosing a more harmful option for convenience, profit, or status SHOULD trigger CSV_REDESIGN_REQUIRED or a high residual RLS penalty, depending on severity.

9.3.5 Necessity Corridor

A harmful option MAY receive CSV_PASS_WITH_CONTROLS or CSV_EMERGENCY_PROVISIONAL only when it satisfies the necessity corridor:

Corridor condition	Requirement
Necessity	The option serves a real need, not mere convenience.
No better feasible alternative now	Cleaner or safer options are unavailable, unaffordable, disproportionate, or not timely under the declared reality surface.
Boundedness	Harm is limited in time, space, exposure, scale, and recurrence.
Mitigation	Controls reduce harm materially.
Transparency	The harm is not hidden, greenwashed, or falsely described.
Reversibility / exit	The option does not create avoidable lock-in.
Transition	A path to a cleaner, safer, or more viable option is declared where material.
Review	Monitoring and re-evaluation triggers are declared.

The necessity corridor prevents CSV from becoming impossible purity ethics while preserving its ability to block structural laundering.

9.4 CSV Diagnostic Concern Scale (Normative guidance)

CSV MAY use internal diagnostic concern levels. These levels are not an RLS score and MUST NOT be treated as final welfare ranking.

Level	Meaning	Default disposition
0	No material CSV concern.	CSV_PASS.
1	Minor residual concern.	Carry to RLS and monitoring.
2	Material concern requiring controls.	CSV_PASS_WITH_CONTROLS if controls are binding.
3	Serious concern requiring redesign or higher review.	CSV_REDESIGN_REQUIRED or CSV_ESCALATE.
4	Gate-failing concern.	CSV_FAIL.

Suggested diagnostic dimensions:

- Containment integrity
- Structural viability
- Hollowing risk
- Dependency or lock-in risk
- Substitution pressure
- Mitigation adequacy
- Monitoring adequacy
- Reversibility / exit
- Accumulation risk
- Governance or legitimacy stress

Default status rule:

If any material diagnostic = 4, CSV_FAIL. If any material diagnostic = 3, CSV_REDESIGN_REQUIRED or CSV_ESCALATE. If any material diagnostic = 2, CSV_PASS_WITH_CONTROLS only if controls are binding and auditable. If all material diagnostics are 0 or 1, CSV_PASS and residuals carry to RLS.

9.5 Formal CSV Predicate (Normative)

Let A_{adm} be the set of options that pass NCRC and TRC, subject to Reality Grounding supporting the claim boundary.

For each option a in A_{adm} , compute $CSV_status(a)$.

Default selectable set:

$A_{sel} = \{a \text{ in } A_{adm} : CSV_status(a) \text{ in } \{CSV_PASS, CSV_PASS_WITH_CONTROLS, CSV_NOT_MATERIAL\}\}$

Emergency provisional options are not ordinary members of A_{sel} . A CSV_EMERGENCY_PROVISIONAL option may be temporarily selected only through emergency protocol, with explicit time limit, harm cap, mitigation, review, and no full alignment claim.

CSV_NOT_MATERIAL is pass-equivalent only for the declared claim boundary after the PCC records the rationale that no material CSV burden is present. It MUST NOT be used to waive Structural Viability subchecks when execution feasibility, resource closure, dependency closure, reversibility, operational capacity, or internal coherence is material. In high-stakes or execution-bearing runs, CSV_NOT_MATERIAL may apply to the containment-integrity dimension while Structural Viability remains separately required where execution feasibility is material. If later evidence makes CSV material, CSV MUST be rerun.

RLS is applied only to A_{sel} , except where emergency protocol explicitly invokes fallback comparison.

9.6 UCI, HOI, and Structural Diagnostic Placement (Normative)

CSV is the authority layer for cross-scale structural integrity. UCI, HOI, Adaptive Reserve, and related structural indicators are diagnostic instruments. When such indicators are material to whether an option degrades a containing union, hollows out capacity, creates dependency, or undermines structural viability, they are evaluated inside CSV before RLS.

UCI/HOI are not public cascade stages. They are not separate gates by default. They are structural diagnostics inside CSV when material, and residual tie-break or monitoring signals only after options are already selectable and RLS is tied, close, uncertainty-overlapped, or non-decisive.

If UCI/HOI or equivalent structural diagnostic evidence is required for a CSV determination and unavailable, the implementation SHALL record CSV_STRUCTURAL_DIAGNOSTIC_UNAVAILABLE or the more specific CSV_UCI_HOI_UNAVAILABLE flag. A run with unresolved material diagnostic unavailability SHALL NOT claim an unqualified CSV pass or Tier 3 CSV conformance.

9.7 CSV Fail Conditions (Normative)

CSV MUST fail, redesign, or escalate an option when any of the following are material and unresolved:

1. Container destruction: the option damages the system that makes the claimed benefit possible.
2. Forced externalization: the option works by pushing burden onto people, ecosystems, or future agents who cannot consent, contest, or recover.
3. Structural non-viability: the option cannot execute under real resource, dependency, capacity, authority, or contradiction constraints.
4. Hollowing: the option improves surface metrics while weakening the deeper capacity the metric was meant to serve.
5. Lock-in or dependency: the option creates avoidable dependency that makes better futures harder.
6. Unbounded accumulation: individually small harms scale into systemic degradation when repeated.

7. Unmonitored invisible harm: material harms exist without monitoring, audit, review, or feedback path.
8. Better feasible alternative ignored: a much cleaner or safer feasible option exists, but the harmful option is chosen for convenience, status, or profit.

9.7A Distributional Burden and Non-Erasure Review (Normative)

When Section 2.4A is triggered, CSV MUST create a DistributionalBurdenReview before the selectable set is formed. The review is a record feeding the existing CSV state machine; it does not create a new gate or a new pass-equivalent CSV verdict.

The review SHALL evaluate severity, incidence, concentration, duration, recurrence, reversibility, remedy, ability to consent/contest/exit/recover, forced externalization, invisible or unmonitored burden, constitutive mitigation, accountable carrier, monitoring, and materially less-harmful feasible alternatives.

CSV_PASS is available only when the burden is bounded within the declared tolerance and no constitutive controls are required. CSV_PASS_WITH_CONTROLS is available only for the controlled option. CSV_REDESIGN_REQUIRED applies where a materially less-harmful feasible alternative or redesign exists. CSV_ESCALATE applies where evidence or authority is insufficient. CSV_FAIL applies to uncontained, beyond-tolerance, forced, invisible, unrecoverable, or structurally degrading burden.

A population mean MUST NOT be the sole basis for clearing a severe concentrated burden. If every feasible option imposes severe unavoidable burden, the run SHALL identify the least-harm selectable option, make mitigation and remedy constitutive, preserve representation and appeal, record residual-harm authorization separately, and maintain transition/review triggers.

CSV_NOT_MATERIAL is unavailable where any of the following is plausibly material: irreversibility; severe third-party externality; uncontrolled dependency; rights-adjacent exposure; material lock-in; unbounded accumulation; inability to consent, contest, exit, or recover; or a challenger-supported distributional omission. A contrary disposition requires independent reviewer rationale.

9.8 CSV Pass Conditions (Normative)

CSV MAY pass an option when harms are visible, rights floors are intact, tail risk is bounded, containment burden is limited, structural viability is real, externalities are mitigated, residual harms are carried into RLS, monitoring exists where needed, lock-in is avoided or justified, better alternatives are unavailable or disproportionate, and transition or review plans exist where material.

Plain-language test:

Harm does not equal fail. Hidden, unbounded, unjust, irreversible, or structurally degrading harm equals fail, redesign, escalation, or emergency provisional treatment.

9.8A Binding-control and obligation-integrity condition (Normative)

A CSV_PASS_WITH_CONTROLS is claimable only when every selectability-material control satisfies Section 4.10A. Controls that are promised, unowned, under-resourced, unmonitored, consequence-free, dependent on silent carrier compensation, or alterable without governed change control SHALL NOT upgrade CSV status. If carrier nonperformance, control expiry, or unauthorized control modification would leave a material harm uncontained, the option must be redesigned, escalated, suspended, or failed according to severity and available alternatives.

9.8B Human-compensation-load diagnostic (Normative when triggered)

CSV SHALL test hidden human compensation when current performance, proposed automation, scaling, outsourcing, or role removal could conceal a dependency on undocumented judgment or recovery work. A bounded and supported human contribution may remain part of a viable design. A material, unrecorded, involuntary, saturated, single-carrier, or non-substitutable compensation burden requires binding controls,

capacity restoration, documentation, redistribution, redesign, escalation, or failure. This diagnostic remains inside CSV and does not create a new gate or presume that all human adaptation is undesirable.

9.9 Worked Normative Example: Backup Diesel Generator (Informative but Canon-Guiding)

A small organic cafe may need a backup diesel generator during power cuts. The generator produces smoke, noise, local air pollution, and emissions. Without it, the cafe may lose food, income, jobs, community function, and operational continuity.

CSV MUST NOT treat this as an automatic fail merely because harm exists. CSV MUST also not allow the cafe's local survival benefit to hide uncontained harm.

Illustrative dispositions:

Generator condition	CSV disposition	Rationale
Generator used indoors, near air intakes, or near vulnerable people without controls.	CSV_FAIL or NCRC escalation.	Serious health/safety exposure and containment breach.
Outdoor generator, frequent use, no mitigation, no review, no transition.	CSV_REDESIGN_REQUIRED.	Harm is visible but insufficiently contained.
Outdoor generator, limited outages, exhaust away from people, maintenance, fuel controls, noise limits, and transition plan.	CSV_PASS_WITH_CONTROLS.	Harm is bounded, necessary, mitigated, transparent, and reviewable.
Battery/inverter feasible now.	CSV_PASS; likely higher RLS.	Lower containment burden.
Solar + battery + emergency diesel only.	CSV_PASS; likely stronger RLS if feasible.	Better resilience/environment balance.
Closing during outages.	CSV_PASS but RLS context-dependent.	Avoids smoke but may harm livelihood, waste food, and reduce community function.

Canonical lesson:

CSV is not “diesel bad, therefore fail.” CSV is “diesel carries a containment burden; if necessary, bounded, mitigated, transparent, monitored, and transitional, it may conditionally pass and be ranked against cleaner options by RLS.”

9.9A Qualified Controls, Timing and Recovery (Normative)

A qualified control envelope is the existing control/obligation record made explicit: control and option identity, hazard/endpoint, trigger, scope, accountable authority, operational carrier, capacity and resources, response-time bounds, effectiveness evidence and uncertainty, dependencies, failure modes, reversibility, residual harms, expiry, fallback and outcome verification. Where a field already exists, reference it rather than create a competing status object. The control is constitutive of the option; crediting it changes the option and requires the affected prior checks.

Runtime prevention credit requires evidence that the control can reach the required safe state before the earliest unacceptable endpoint under the declared credible timing bounds. A useful conservative test is $T_{\text{control_upper}} + \text{margin} < T_{\text{harm_lower}}$. The margin and bounds are governed domain inputs, not universal constants. Failure of this test forbids claiming timely prevention of that endpoint; it does not imply that all later mitigation, recovery or independent ex-ante protection has zero value. Credit those different endpoints separately when supported.

Effectiveness η in $[0,1]$ and residual transmission $\psi = 1 - \eta$ are permitted explanatory parameters only for a specified controllable pathway with adequate warrant. A generic equation, boolean shield flag, cryptographic

key or high shock velocity cannot reduce real loss. Unknown effectiveness receives no asserted protective discount for the decision-critical endpoint. Reserve benefits require units, accessible capacity, consumption/depletion and replenishment evidence. Do not multiply multiple controls as independent without support, conceal costs or externalized burdens, or let private reserves determine entitlement to protection.

Observability, available intervention capacity and authority SHALL be distinguished. Adequate observation and control are task-, endpoint-, horizon- and configuration-specific; global full-state observability or controllability is not a universal prerequisite for a particular safe stop. Monitoring alone is not control, and a stop button alone does not prove a feasible timely safe state.

Recovery proposals must pass applicable RG, RF/NCRC, TRC and CSV, with separately bound authority and execution controls. Maintain observation, proposed action, authorization, attempt and verified outcome as separate records. Repeated delivery, retries, restarts or failover must not multiply an authorized real-world intervention. Unknown execution outcome requires reconciliation rather than blind retry; an opposite-signed score is not an inverse physical intervention.

9.10 RLS Relationship (Normative)

RLS ranks only options that survive RG claim support, the Rights Floor (formal NCRC), TRC, and CSV. Decision-material residual welfare consequences identified during CSV SHOULD be represented once through their ordinary evidenced EffectTokens and RLS cells, including pollution, dependency, lock-in, ecological burden, social trust loss and meaning/agency degradation. Do not add a separate CSV penalty for a consequence already included in the residual field; failed qualification remains outside ordinary RLS.

RLS MUST NOT rescue CSV failure. CSV protects the doorway into ranking; RLS ranks what may legitimately enter.

9.11 CSV Parameter Governance and Escalation (Normative)

CSV thresholds, diagnostic concern levels, necessity-corridor judgments, UCI/HOI trigger thresholds, and structural-viability sufficiency standards MUST be recorded in the PCC when material. Tier 3 runs SHOULD include sensitivity analysis over near-binding CSV determinations.

If no option passes CSV, the run MUST NOT pretend a good option exists. Required actions include option redesign, additional evidence gathering, transition planning, escalation to a higher governance tier, emergency provisional protocol if genuinely necessary, or refusal of deterministic selection.

CSV does not override NCRC or TRC. It operates after admissibility and before residual welfare ranking.

SECTION 10: OPTIMIZATION LAYER: RIPPLELOGIC SCORE (RLS)

10.0 Purpose and Scope (Normative)

RLS ranks selectable options by expected welfare improvement across unions and dimensions, after ripple propagation and saturation, under declared weights and applicability mask rules.

RLS MUST NEVER override NCRC, TRC, or CSV. It operates only on selectable options A_sel.

10.0A RLS Claim Boundary and Profile-First Reporting (Normative)

RLS ranks selectable options through a bounded, governance-defined, constitution-conditional residual welfare-impact index under the declared evidence, effect representation, reference structure, baseline, probability semantics, temporal treatment, welfare-inclusion hypotheses, uncertainty model, applicability mask, interaction assumptions, and weight constitution.

RLS does not establish a natural quantity of total welfare, interpersonal utility, universal cross-substrate value, risk-adjusted expected utility, moral truth, or an unconditional ethically best option. A common numerical range is not proof of a common cardinal welfare unit.

Profile-first rule. Before the scalar summary, claim-bearing reports SHALL show the 7×7 profile, active and non-maskable cells, subgroup and lower-tail findings, worst materially affected groups, uncertainty and dependence posture, evidence maturity, reference/baseline family, weight-profile status, material sensitivities, and any incomparability or non-decisiveness.

The point-score leader, framework-selected option, authority-selected option, and execution-authorized action are four distinct states and MUST NOT be collapsed.

10.1 Normalized RLS Definition (Normative)

Define the effective residual cell weight:

$$q(u,d) = w_u \cdot v_d \cdot m(u,d) \cdot \kappa(u,d)$$

Weight-separability assumption (Normative disclosure). The reference construction factorizes scope prominence and dimension prominence through $w_u \cdot v_d$, with $\kappa(u,d)$ available only as a bounded, declared interaction correction. This is a transparent model assumption, not an empirically validated claim that scope and dimension importance are universally separable. Tier 3 and decision-material or high-stakes Tier 2 runs MUST test a plausible non-separable cell-weight matrix or bounded interaction alternative whenever scope-by-dimension interaction could materially change ranking or decisiveness. If the result changes, set the applicable sensitivity/review state and treat deterministic selection as unresolved until governed review is completed.

and the active effective mass:

$$Q = \sum_u \sum_d q(u,d)$$

For $Q > 0$:

$$RLS(a) = [\sum_u \sum_d q(u,d) \cdot I_{prop_welfare}(u,d,a)] / Q$$

If $Q = 0$, RLS is undefined. Return `RLS_NO_ACTIVE_MASS`, repair the mask/weight declaration, narrow the claim, or refuse deterministic ranking; zero MUST NOT be imputed.

Constructed-index caution (Normative claim boundary). RLS is a bounded, governance-defined index over declared, baseline-relative welfare-impact judgments. It is not a natural welfare unit, an interpersonal utility meter, empirical validation of the 7×7 structure, or a scalar proof of moral truth. Its arithmetic is authoritative only within the declared model and selectable set, with evidence status, uncertainty, weighting authority, subgroup visibility, and sensitivity disclosed.

Where:

- w_u are union weights (HDW-governed; Section 13; Section 13.1 constitutional floors controlling)
- v_d are dimension weights (HDW-governed; Section 13; Section 13.1 constitutional floors controlling)
- $m(u,d)$ is the applicability mask for RLS aggregation only
- $\kappa(u,d)$ is a declared positive cell multiplier, default 1
- $I_{prop_welfare}(u,d,a)$ is the post-propagation, post-saturation welfare impact

Normalization by Q preserves the public interpretation of RLS within $[-1,+1]$ when cell impacts lie within $[-1,+1]$, including when masks or non-default κ values change the active effective mass.

$\kappa(u,d)$ governance rule (Normative). The default is $\kappa(u,d) = 1$. Non-default $\kappa(u,d)$ values are permitted only for pre-declared, governance-approved cell-level emphasis that does not alter admissibility semantics. Unless a stricter charter rule is declared, $\kappa(u,d)$ MUST remain positive and bounded within $[0.5,1.5]$. A non-default $\kappa(u,d)$ MUST NOT reverse sign, erase or dilute rights-covered or catastrophe-relevant harm, or function as a back-door mask, rights-floor rewrite, or prominence override. Tier 3 runs using non-default $\kappa(u,d)$ MUST record justification, authorized source, affected cells, Q , and sensitivity reruns in the PCC.

Local-use note (Informative). Section 10.2A defines a Practical Local Scope Scoring profile. PLSS preserves the same normalized equation and constructs w_u within the residual mass above constitutional floors.

10.1A Exact Contribution Decomposition and Notation (Normative clarification)

Using the canonical effective cell weight $q(u,d)=w_u v_d m(u,d) \kappa(u,d)$ and $Q=\sum_{(u,d)} q(u,d)$, define $C_{ud}(a)=q(u,d) I_{\text{prop,welfare}}(u,d,a)/Q$ for $Q>0$. Then $RLS(a)=\sum_{(u,d)} C_{ud}(a)=\sum_u C_u(a)=\sum_d C_d(a)$, where $C_u=\sum_d C_{ud}$ and $C_d=\sum_u C_{ud}$. These are alternate decompositions of the same score; adding the row and column totals together would double count it.

The Contribution Matrix SHALL be available for a reported numerical RLS, together with the active-cell basis, unrounded values, rounding policy and the governing weight profile. Report gross positive and gross negative final-cell contributions where useful; their sum is RLS, but they do not recover harms hidden by within-cell netting. Preserve effect-level and subgroup findings separately. A display of 100 times RLS is a rescaled index, not a percentage of alignment, safety, welfare delivered or ethical correctness.

The cell decomposition is exact algebra; causal-source attribution through tanh, interaction models or other nonlinear transforms is not automatically additive. Report source contributions before saturation or use a named, assumption-disclosed attribution method. Do not present invented causal percentages. UNKNOWN is not a number and cannot be repaired by an attractive color, clipping or a zero-filled chart.

K denotes the propagation matrix when used as a matrix; $K(u,d,a)$ remains the existing set-of-impact-instances notation where explicitly indexed as a set. The probability-ownership and scenario-first rules, confidence sensitivities, rights/base versus welfare streams, normalization, every-contender discrimination and authority separation are unchanged.

10.2 Applicability Mask Rules (Normative)

Masking exists to avoid meaningless aggregation in some contexts, but is a primary gaming vector. Therefore:

Definition (Normative). Masking means setting $m(u,d)=0$, which affects RLS aggregation only.

Omission/exclusion means failing to compute, disclose, or evaluate a required cell for NCRC, TRC, or CSV.

Masking SHALL NOT be used as a substitute for omission rules.

Masking affects RLS aggregation only. Whether a cell must be constructed is determined by active-cell rules (Section 10.3) and tier coverage requirements, not by masking. Accordingly, masking MUST NOT be used to skip or modify:

- Impact construction (Section 5),
- Propagation (Section 6),
- NCRC checks (Section 7),
- TRC loss and CVaR computation (Section 8),
- Or containment evaluation (Section 9).

Non-maskable cells (MUST be unmasked):

- All rights-covered cells in any C_r
- All catastrophe cells in C_{cat}

Non-maskable persistence rule (Normative). Rights-covered cells remain non-maskable even after NCRC admissibility filtering. This prevents post-gate dilution and preserves rights-salience when ranking admissible options, especially under adversarial pressure and Goodhart dynamics.

Rationale (Informative). The admissibility gates prevent selecting options that violate non-compensatory rights floors, but they do not by themselves prevent near-miss normalization among admissible options. Keeping rights and catastrophe cells in RLS aggregation preserves continuous optimization pressure toward larger safety margins, reduces proxy gaming pressure (for example, satisfying a threshold while degrading adjacent rights conditions), and makes rights and tail-risk performance legible for governance review and accountability.

Baseline consistency for non-maskable gate cells (Normative). For any cell in C_r or C_{cat} that remains included in RLS under the non-maskable persistence rule, the welfare-stream baseline SHALL also use the gate floor-reference baseline unless a governed dual-baseline method is explicitly declared. Where a dual-baseline method is declared, the floor-reference value SHALL remain the controlling value for gate-consistency interpretation and validator conformance.

Nonlinear baseline-sensitivity boundary (Normative clarification). Because cell impacts are saturated nonlinearly, shifting from STATUS_QUO to FLOOR_REFERENCE is not guaranteed to add a rank-neutral constant across options. RLS values for non-maskable cells using FLOOR_REFERENCE therefore describe movement relative to the rights-safe or safe-corridor reference, not ordinary net change from status quo. Tier 3 and high-stakes Tier 2 runs MUST report a governed status-quo counterfactual for ranking whenever a defensible status-quo welfare baseline exists and the baseline choice could alter ranking or decisiveness. If ordering changes, the run is BASELINE_RANK_SENSITIVE; the floor-reference result remains controlling for gate consistency, but deterministic residual selection MUST be narrowed, escalated, or routed to accountable authority rather than presented as baseline-invariant.

- Any cells mandated by minimum governance coverage for the run (declared in PCC; default includes at least U_1 , primary affected unions, and U_7 Environment when environmental relevance exists)

Audit rule: If any non-maskable cell is masked, set audit_flag RIGHTS_CELL_MASKED_INVALID (or CATASTROPHE_CELL_MASKED_INVALID) and PCC is INVALID.

10.2A Practical Local Scope Scoring (PLSS) (Normative)

Purpose. PLSS specifies a canon-consistent local-use profile for RippleLogic Score (RLS) when a decision is primarily local in its direct effects, but broader scopes remain morally relevant as rights, tail-risk, and containment guardrails. PLSS allows routine decisions to concentrate most residual welfare-ranking attention on the scopes that are actually live in the decision, without weakening NCRC, TRC, CSV, non-maskable cell rules, or Section 13 constitutional weight floors.

Interpretive anchor (Informative). Union Scopes are not a closed list of who matters. They are the stable scopes of aggregation and accountability used by the framework. Stakeholder instances feed the scopes; the scopes ensure that impacts do not disappear.

Terminology lock (Informative). LSS names the upstream stakeholder object, PLSS names the residual weight-construction procedure used within RLS, and FLGG - Focus Local, Guard Global - names the governing design principle.

10.2A.0 Attention-Safety Separation Principle (Normative, architecturally grounded). PLSS operates only within the residual welfare-ranking layer. Admissibility under RF/NCRC, TRC, and CSV SHALL remain invariant under any valid PLSS weight construction satisfying Sections 10.2A.4-10.2A.6. Therefore, a change in PLSS prominence declarations or PLSS weight allocation SHALL NOT convert an inadmissible option into an admissible option. This principle follows from the lexicographic separation of admissibility gates from the welfare-ranking layer.

PLSS does not determine who matters; it determines how residual welfare-ranking attention is distributed after universal admissibility conditions have already been imposed.

10.2A.1 Relation to existing gates (Normative)

PLSS applies only to welfare ranking within Section 10. It SHALL NOT alter:

- NCRC evaluation over the required rights coverage sets C_r ,

- TRC evaluation over the catastrophe cell set C_{cat} ,
- Containment requirements under Section 9,
- non-maskable cell requirements under Section 10.2,
- constitutional weight floors defined in Section 13.1.

A PLSS run therefore remains:

- full-scope for admissibility, where required by RF/NCRC, TRC, and CSV,
- local-scope only in residual welfare emphasis, after admissibility has been preserved.

Structural diagnostic preservation under PLSS (Normative). PLSS SHALL NOT down-weight, omit, mask, or defer structural diagnostic evidence required for Containment. If a local-scope run uses PLSS but plausible containment-relevant UCI/HOI, Adaptive Reserve, lock-in, hollowing, governance-spillover, ecological-spillover, or institutional-capacity evidence exists outside the prominent scopes, that evidence remains admissibility-relevant. The run SHALL escalate, broaden scrutiny, or refuse the stronger local-scope claim when containment-relevant structural evidence is unavailable or unresolved.

10.2A.2 Prominent scopes and guardrail scopes (Normative)

For a declared run, each operational Union Scope $u \in \{U1, \dots, U7\}$ may be classified in the PCC as one of the following:

- Prominent scope: a scope whose direct or near-field ripple is substantial enough that it should receive meaningful residual welfare weight in the current run.
- Guardrail scope: a scope whose direct ripple in the current run is limited, but which remains protected by constitutional floor weight and by admissibility constraints.

Guardrail scope classification SHALL NOT be interpreted as moral irrelevance. A guardrail scope may carry little or no residual weight in a local run, but it remains:

- constitutionally represented through Section 13 floors,
- admissibility-relevant where rights, catastrophe, or containment conditions apply.

10.2A.3 Representation rule: instances first, scopes second (Normative)

PLSS SHALL preserve the Section 2.2A and Section 5 representation architecture:

- stakeholder discovery is performed at the instance level,
- instances are mapped to one or more Union Scopes,
- impacts are aggregated within each scope row,
- scope rows are then combined by RLS.

Let I_u denote the declared stakeholder instance set mapped to scope u .

For non-rights scoring, the default aggregation rule remains the canonical population-weighted mean across affected instances:

$$I_{agg}(u,d,a) = [\sum_{i \in I_u} \pi_i \cdot I_i(u,d,a)] / [\sum_{i \in I_u} \pi_i]$$

where π_i is the declared aggregation weight for instance i , for example population share, exposure share, or another declared ReachBasis-consistent basis.

For rights-covered cells, aggregation SHALL NOT average away the worst-off subgroup. Rights evaluation remains governed by the worst-off subgroup operator:

$$I_{rights}(u,d,a) = \min_{g \in G(u,d)} I_g(u,d,a)$$

This subsection does not change Section 7 semantics.

10.2A.4 Local prominence signal and residual allocation rule (Normative)

PLSS introduces a declared local prominence signal q_u for each scope u , used only to distribute the residual union-weight mass that remains after constitutional floors.

Each q_u SHALL be recorded in the PCC and SHALL be justified by reference to the current decision context. A recommended default construction is a declared combination of:

- expected magnitude of ripple in scope u ,
- probability that the effect materializes,
- temporal persistence or duration,
- propagation potential beyond immediate contact.

A recommended default operator is:

$$q_u = \psi(M_u, P_u, T_u, G_u)$$

with $M_u, P_u, T_u, G_u \in [0,1]$, and a default transparent choice:

$$\psi(M_u, P_u, T_u, G_u) = (M_u P_u T_u G_u)^{(1/4)}$$

This recommended operator is informative by default unless separately elevated by governance rule; the required part is that the PCC declare the operator used.

Prominence evidence discipline (Normative). Prominence sub-component declarations for q_u , including magnitude, probability, temporal persistence, and propagation potential, are subject to the same PCC justification and evidence documentation discipline that applies to impact-instance parameters under Section 5.2, including ReachBasis declaration where applicable. Tier 2+ runs SHALL record the evidentiary basis used to justify each declared prominence component.

Cross-deployment comparability note (Normative guidance). The geometric-mean operator remains the default comparability anchor for q_u . Any deployment using a non-default operator MUST document why the operator preserves anti-gaming discipline, boundedness, and monotonicity, and SHOULD include a sensitivity comparison against the geometric-mean default in the PCC or companion validation record.

Define residual shares:

$$\rho_u = q_u / \sum_j q_j, \text{ if } \sum_j q_j > 0$$

If all $q_u = 0$, the residual SHALL be allocated uniformly unless another declared fallback is specified in the PCC.

Let w_u^{floor} be the Section 13.1 constitutional floor for scope u . Then the PLSS union weights are:

$$w_u^{\text{PLSS}} = w_u^{\text{floor}} + (1 - \sum_j w_j^{\text{floor}}) \rho_u$$

subject to the canonical constraints:

$$w_u^{\text{PLSS}} \geq 0$$

$$\sum_u w_u^{\text{PLSS}} = 1$$

$$w_u^{\text{PLSS}} \geq w_u^{\text{floor}}$$

for every scope u .

Interpretation (Informative). PLSS does not re-write the constitution. It lets local runs place most of the residual emphasis on nearby scopes while retaining the constitutional baseline presence of broader scopes.

10.2A.5 PLSS RLS equation (Normative)

For any selectable option $a \in A_{\text{sel}}$, the local-scope welfare ranking is:

$$q_{\text{PLSS}}(u,d) = w_u^{\text{PLSS}} \cdot v_d \cdot m(u,d) \cdot \kappa(u,d)$$

$$Q_{\text{PLSS}} = \sum_u \sum_d q_{\text{PLSS}}(u,d)$$

$$\text{RLS_PLSS}(a) = [\sum_u \sum_d q_{\text{PLSS}}(u,d) \cdot I_{\text{prop_welfare}}(u,d,a)] / Q_{\text{PLSS}}, \text{ for } Q_{\text{PLSS}} > 0$$

where:

- w_u^{PLSS} is defined by Section 10.2A.4,
- v_d remains governed by Section 13,
- $m(u,d)$ remains governed by Section 10.2,
- $\kappa(u,d)$ remains governed by Section 10.1,
- $I_{\text{prop_welfare}}(u,d,a)$ remains governed by Sections 5–6.

If no kernel propagation is declared, then:

$$I_{\text{prop_welfare}}(u,d,a) = I_{\text{dir}}(u,d,a)$$

per the existing NONE propagation pathway.

10.2A.5A Discrimination threshold and non-decisive handling (Normative). The discrimination threshold defined in Section 10.4 applies identically to RLS_PLSS . If the lead against any other selectable option is non-decisive under the declared uncertainty method, the tie-break chain in Section 11.5 MUST be applied. PLSS SHALL NOT bypass uncertainty handling, decisiveness testing, or tie-break governance.

10.2A.5B PLSS prominence sensitivity (Normative)

When PLSS is used in Tier 3, the PCC MUST rerun the welfare ranking under at least two alternative valid prominence constructions:

- (i) the uniform residual fallback, and
- (ii) a broadened guardrail-sensitive vector in which every scope classified Guardrail receives q_u at least 0.10 times the maximum q_u among prominent scopes, unless a stricter charter rule is declared.

If the ranking order of selectable options or the selected option changes under these alternative valid prominence constructions, the run MUST set audit flag `PLSS_PROMINENCE_SENSITIVE` and document the sensitivity in the PCC.

When `OTHER_DECLARED` uses a direct-q normalization such as `DIRECT_Q_NORMALIZATION_V1`, the published sensitivity bundle MUST show at minimum: the declared direct-q result; the default geometric-mean prominence result; the uniform-residual fallback; and a broadened guardrail-sensitive result where plausible wider-scope effects are increased. The PCC MUST state whether ranking, decisiveness, or authority selection changes under each construction.

This requirement does not alter admissibility, which remains governed by the Attention-Safety Separation Principle and the upstream gates.

10.2A.6 Masking and non-maskable rule under PLSS (Normative)

PLSS SHALL NOT modify the Section 10.2 mask rule.

In particular:

- masking remains RLS aggregation only,
- masking SHALL NOT be used as omission,

- rights-covered cells and catastrophe-relevant cells remain non-maskable where required by the canon,
- reduced-scope or local-scope presentation SHALL NOT bypass admissibility on required cells.

Thus PLSS narrows practical emphasis, not moral coverage.

10.2A.7 Escalation triggers (Normative)

A run SHALL NOT remain in PLSS-only posture, or SHALL be rerun with broader scrutiny, if any of the following are plausible:

1. rights-covered harm beyond the prominent scopes,
2. catastrophe relevance or irreversible loss,
3. meaningful externalities crossing Community, Organization, Polity, Humanity / Global Coordination, or Biosphere pathways,
4. lock-in, precedent, or nontrivial governance spillover,
5. a credible challenger claim that a stakeholder class or scope was omitted,
6. Containment requires unavailable evidence for containing scopes,
7. repeated comparable runs continue to use local-scope framing without the required next-run expansion plan.

Audit flag (Normative). If one or more escalation triggers in Section 10.2A.7 is plausibly active, but the run proceeds in PLSS posture without escalation or documented justification, implementations MUST set audit flag `PLSS_ESCALATION_TRIGGER_UNRESOLVED` with severity `ESCALATE`. The PCC MUST then either record the broader posture adopted or provide a specific justification for non-escalation, including evidence basis and reviewer sign-off where required by tier. If no such disposition is recorded, the validator MUST reject the run as non-conformant.

Boundary heuristic note (Normative guidance). Implementations MAY use lightweight pre-screen heuristics to decide whether FLGG escalation is required before full local scoring proceeds, for example trigger scans over rights exposure, catastrophe relevance, cross-scope externality breadth, lock-in, precedent, and containment-evidence availability. Any such heuristic is advisory only: once a trigger is plausibly active, the canonical escalation requirements of Section 10.2A.7 control.

Ex ante method declaration rule (Normative). Any method choice that can materially affect decisiveness, ranking within the selectable set, or tie-break outcome, including PLSS prominence construction, alternative PLSS operator, scope-posture escalation disposition, and UCI comparison method, SHOULD be declared ex ante in the PCC where practicable. If first declared only after comparative outcomes are observed, the PCC MUST record a reason code, reviewer sign-off where required by tier, and a sensitivity note sufficient for independent review.

Scope posture classification note (Normative). When PLSS is used, each scope SHALL be classified in the PCC as either Prominent or Guardrail. v9.0 removes the undefined “Mixed” category to reduce audit drift; if a deployment wants finer descriptive language, it MAY add non-canonical comments, but the PCC classification field SHALL remain binary.

10.2A.8 Plain-language reading rule (Normative)

When not to use PLSS (Normative guidance statement). PLSS is intended primarily for local or moderately bounded decisions in which direct welfare effects are concentrated and no escalation trigger is plausibly active at the outset. PLSS SHALL NOT be used as the primary posture for irreversible cross-border policy, large-scale public AI deployment, biosecurity governance, active military or security operations, high-lock-in infrastructure decisions, or any case in which ET-1 through ET-7 is plausibly active at decision entry. Such cases require a broader deliberative posture or higher-tier structured evaluation.

For instructional and public-facing use, the framework MAY describe PLSS with the following plain-language sentence:

“In local decisions, score most strongly where the ripple is near, but keep every wider scope present through floors and guardrails. PLSS narrows residual attention, not moral coverage. It MAY change welfare ranking within the selectable set, but it MUST NOT make an inadmissible option admissible.”

This sentence is interpretive only and SHALL NOT override the equations above.

10.2A.9 Worked example: family café local run (Informative)

Decision. A family café must choose between:

- Option A: cheap single-use cups,
- Option B: reusable deposit-return cups with a local wash partner,
- Option C: ultra-cheap cups from an unverified supplier with plausible toxic-material or labor-abuse exposure.

Declared stakeholder instances by scope:

- U1 Self: 1 instance
- U2 Household: 1 instance
- U3 Community: 5 instances
- U4 Organization: 2 instances
- U5 Polity: 3 instances
- U6 Humanity / Global Coordination: 1 instance
- U7 Biosphere: 1 instance

This run illustrates the distinction between instances and scopes. The café does not score Community in the abstract. It scores the affected Community instances, then aggregates them into the U3 row.

Declared local prominence signals:

$$q = (U1, U2, U3, U4, U5, U6, U7) = (0.95, 0.85, 0.55, 0.50, 0.10, 0.03, 0.02)$$

These sum to 3.00.

Under Section 13.1, the union floors are:

$$w^{\text{floor}} = (0.20, 0.06, 0.06, 0.06, 0.08, 0.10, 0.10)$$

The floor total is 0.66, so the residual is 0.34. The resulting PLSS weights are:

$$w^{\text{PLSS}} \approx (0.308, 0.156, 0.122, 0.117, 0.091, 0.103, 0.102)$$

Interpretation. Self, Household, Community, and Organization receive most of the residual attention. Polity, Humanity / Global Coordination, and Biosphere remain meaningfully present because the current constitution keeps them above zero even in local runs.

Table 10-1. Illustrative local prominence signals and PLSS union weights

Union Scope	q_u	w_u^{floor}	w_u^{PLSS}
U1 Self	0.95	0.20	0.308
U2 Household	0.85	0.06	0.156
U3 Community	0.5	0.06	0.122

Union Scope	q_u	w_u^{floor}	w_u^{PLSS}
	5		
U4 Organization	0.5 0	0.06	0.117
U5 Polity	0.1 0	0.08	0.091
U6 Humanity / Global Coordination	0.0 3	0.10	0.103
U7 Biosphere	0.0 2	0.10	0.102

Illustrative scope-level welfare composites:

Table 10-2. Illustrative scope composites by option

Scope	$Q_u(A)$	$Q_u(B)$
U1 Self	0.25	0.18
U2 Household	0.10	0.35
U3 Community	-0.06	0.41
U4 Organization	0.20	0.32
U5 Polity	-0.10	0.12
U6 Humanity / Global Coordination	-0.08	0.05
U7 Biosphere	-0.25	0.30

Then:

$$RLS_PLSS(A) \approx 0.0656$$

$$RLS_PLSS(B) \approx 0.2444$$

So Option B wins on welfare ranking among selectable options.

Worked-example uncertainty note. The raw RLS_PLSS values shown in this example do not by themselves determine final selectability. Under the declared uncertainty method, the lead against every other selectable option SHALL be tested against the Section 10.4 discrimination threshold. If the result is non-decisive, the Section 11.5 tie-break chain applies. This example is intended to illustrate PLSS weight construction and score propagation, not to waive the framework's general decisiveness discipline.

Admissibility note. If Option C carries plausible rights-floor or catastrophe concerns through toxic exposure or abusive supply-chain conditions, it may fail NCRC or TRC before RLS ranking. If so, it is not saved by local convenience or low cost.

Lesson (Informative). PLSS does not say "ignore broader scopes." It says: let local decisions place most additional attention where the ripple is near, while still preventing local gain from laundering wider harm.

10.2B Comparison Mask and Active-Cell Symmetry (Normative)

The applicability mask and comparison active-cell set are run-level comparison declarations. The same $m(u,d)$ and comparison-cell basis SHALL be applied across all options being ranked.

Parameter-lock cross-reference. The comparison-mask declaration is fixed before outcome-sensitive ranking under the existing parameter-lock and methodological-integrity rules. A legitimate revision is recorded through governed change control and requires affected qualification, scoring and robustness checks to be rerun; it is not an outcome-selected exception.

A genuine structural difference in applicability MAY be represented by an option-specific exception only when the cell construct is literally undefined or absent for that option, the reason is independent of the desired result, the exception is challenger-visible, and the ranking is stress-tested under a common comparison basis.

Unjustified option-specific masking, or option-specific denominator Q manipulation, sets MASK_ASYMMETRY_INVALID and invalidates the comparative RLS claim. Gate-critical and non-maskable cells remain active regardless of residual masking.

10.3 Uncertainty Handling (Tier 2+ Required for Active Cells) (Normative)

RippleLogic recognizes epistemic uncertainty. Tier 2 and Tier 3 runs MUST record uncertainty and apply the discrimination rule for all active cells (definition below). Tier 1 runs SHOULD record uncertainty when feasible, especially on rights-covered and catastrophe-relevant cells.

Active cell definition (Normative; v9.0 tightened). A cell (u,d) is active for a run if any of the following holds: (i) $m(u,d)=1$ and the cell contributes to RLS aggregation, (ii) the cell lies in any rights coverage set C_r used for NCRC evaluation (regardless of mask status), (iii) the cell lies in the catastrophe cell set C_{cat} used for TRC evaluation (regardless of mask status), (iv) the cell is explicitly populated in the run for any gate, diagnostic, or justification, or (v) the cell is required by the run's minimum governance coverage rules declared in the PCC.

Active cells MUST NOT be treated as unknown-by-default without disclosure. Masking ($m(u,d)=0$) excludes a cell from RLS aggregation only; it does not remove the cell from active status if any of conditions (ii) through (v) hold.

10.3.1 Cell-level uncertainty proxy declarations (Normative)

Definition of $\sigma(u,d,a)$ (normative minimum): PCC MUST declare one of the following cell-level uncertainty proxies and apply it consistently across options:

The PCC MUST choose and record the uncertainty proxy before outcome-sensitive ranking. Method A uses a declared interval half-width; Method B uses the confidence-derived procedure in §10.3.2; Method C is permitted only within the demonstration and calibration boundaries stated below. The selected method, assumptions, validity limits, dependence treatment, and recheck trigger apply consistently across options.

Method A (interval half-width): If the PCC records a confidence interval $[L,U]$ for the cell impact $I(u,d,a)$, set $\sigma(u,d,a) = (U - L)/2$.

Method B (confidence-derived contribution-mass rule): Let x_k be each pre-confidence signed instance contribution in the cell after reach, temporal, likelihood, adjustment, and permitted sentence terms but before cell saturation. Define $A_{cell}(u,d,a) = \min(1, \sum_k |x_k|)$. If the PCC records cell confidence $c(u,d,a) \in [0,1]$, set $\sigma(u,d,a) = (1 - c(u,d,a)) \times A_{cell}(u,d,a)$. This prevents uncertain positive and negative contributions from cancelling to a deceptively small uncertainty value.

Method C (calibrated table): Use a pre-registered mapping from qualitative uncertainty labels (for example, LOW/MED/HIGH) to σ values, stored in the ProofPack or PCC appendix.

10.3.2 Cell-level confidence derivation (Normative; Method B binding)

When the PCC uses Method B (confidence-derived) to compute cell uncertainty for a cell impact estimate, the PCC MUST define cell confidence $c(u,d,a)$ in $[0,1]$ from the set of impact-instance confidences $\{c_k\}$ $k \in K(u,d,a)$ using one declared CellConfidenceAggregationMethod. Implementations MUST apply the same method across all options in the run.

Let the instance contribution weight be defined as:

$$q_k := r_k \cdot \tau(t_k) \cdot \ell_k \cdot e_k \cdot h_k \cdot |\mu_k|$$

and let $Q_{cell} := \sum_{k \in K(u,d,a)} q_k$.

The PCC MUST declare one of the following canonical aggregation methods:

CCAM_MIN_V1 (conservative default):

$$c(u,d,a) := \min_{k \in K(u,d,a)} c_k$$

Use case: high-stakes, high-adversarial-pressure contexts where “weakest-link” confidence is the appropriate posture.

CCAM_WMEAN_V1 (weighted mean): if $Q_{cell} > 0$, $c(u,d,a) := (1/Q_{cell}) \cdot \sum_{k \in K(u,d,a)} (q_k \cdot c_k)$ If $Q_{cell} = 0$,

$$\text{set } c(u,d,a) := \min_{k \in K(u,d,a)} c_k.$$

Edge case (Normative). If $K(u,d,a) = \emptyset$ and the Phantom Instance Rule applies, distinguish point-estimate construction from epistemic uncertainty. The phantom instance uses $c_{point} = 1.0$ only so the declared ignorance penalty is not attenuated in the point estimate. For Method B uncertainty, the cell MUST use $c_{uncertainty} = 0.0$ and therefore $\sigma(u,d,a) = A_{cell}(u,d,a)$ (subject to the ordinary cap), while remaining marked UNKNOWN_IMPACT. An implementation MUST NOT assign zero uncertainty to a phantom-active cell merely because the point-estimate penalty was applied at full strength.

Range rule (Normative). The derived $c(u,d,a)$ MUST be clipped to $[0,1]$ and recorded in PCC for each active cell when Method B is used.

Reconstruction rule. If a worked vector begins from a single-instance post-confidence saturated welfare value $I_{prop_welfare} = \tanh(\beta \cdot c_k \cdot x_{preconf})$, it MUST reconstruct $x_{preconf} =$

$\text{atanh}(I_{prop_welfare})/(\beta \cdot c_k)$, not divide by β alone. For multiple instances, the run MUST retain or reproduce the instance ledger; a cell total generally cannot identify the pre-confidence absolute contribution mass uniquely. **Computability precondition. If the instance ledger cannot be accessed or reproducibly reconstructed, Method B is unavailable unless an explicitly warranted single-instance model uniquely identifies every required term, including nonzero confidence and the applied saturation/propagation path. A net multi-instance cell total MUST NOT substitute for absolute instance contribution mass. Any replacement by Method A or Method C requires its own admissible evidence, a documented PCC reason, and the required requalification and robustness checks; it MUST NOT be chosen merely to obtain a preferred result. If no warranted alternative exists, withhold the uncertainty estimate and refuse a stronger decisiveness claim. An uncalibrated Method C constant remains demonstration-only.**

Scale-consistency rule (Normative). The retained Method B formula is a confidence-derived pre-saturation support proxy, not a calibrated sampling standard deviation or a guaranteed post-saturation error bound. The local derivative of $\tanh(\beta x)$ is $\beta \text{sech}^2(\beta x)$, so the transform can expand small-signal discrepancies or compress large-signal discrepancies. A shrinkage-to-unshrunk difference alone does not establish true statistical uncertainty. Any run at any tier using Method B to support ALLOW_FRAMEWORK_SELECTION or SELECTED_DECISIVE SHALL also declare and pass a transformation-consistent uncertainty treatment on the final I_{prop} scale, including the actual β , scenario aggregation, propagation mode, material dependence and relevant measurement/model uncertainty. One admissible construction is an evidence-supported input interval propagated through the actual nonlinear pipeline and reported as a final-impact interval under Method A; the interval support and interpretation must be disclosed. If a warranted transformed treatment is unavailable, report the legacy proxy and point score only and refuse unique selection. A fixed sigma-times-two test can be informative near zero when $\beta = 2$, but does not by itself establish scale consistency for other β values, kernels or uncertainty models. Historical Method B worked values remain reproducible illustrations, not exemptions from this selection-claim rule.

σ values MUST be clipped to [0,1]. Unless a stricter governed method is declared, $\sigma(u,d,a)$ is interpreted as the pre- κ cell-level uncertainty input and SHALL be aggregated into σ_{RLS} using the same w_u , v_d , $m(u,d)$, and $\kappa(u,d)$ posture used for RLS.

Auditability binding (Normative). The PCC MUST record the source, elicitation basis, or calibration table used for $\sigma(u,d,a)$ and any declared aggregation rule that produces it. If an allowed alternative uncertainty method or aggregation rule would change decisiveness, that sensitivity MUST be disclosed.

Tier 3 uncertainty hardening (Normative). Tier 3 runs MUST rerun decisiveness using $\sigma \times 0.5$ and $\sigma \times 2.0$ stress tests (or the nearest declared equivalent for Method A interval widths) and MUST record whether the every-contender decisiveness result changes.

Direction of the common-multiplier stress. Holding scores and other inputs fixed, reducing positive uncertainty widths raises the absolute Gap and can turn a formerly non-decisive comparison into a decisive one. It cannot overturn an already decisive positive lead solely by increasing its denominator. Therefore sigma-times-0.5 remains a sensitivity diagnostic; it is not universally decision-irrelevant. The exact multiplier relation in Section 10.4 governs the sigma-times-two interpretation.

Cross-tier computed-sensitivity rule (Normative). Whenever a run at any tier actually computes a declared uncertainty, dependence-cluster, baseline, weighting, saturation, or other governed sensitivity and the decisive/non-decisive classification changes under a tested variant, the run MUST set the applicable sensitivity flag, including UNCERTAINTY_DECISIVENESS_SENSITIVE where uncertainty is causal, and SHALL be treated as non-decisive until the tie-break chain or governance review is completed. Optional sensitivity work cannot be ignored merely because the tier did not require it.

Measurement-resolution rule (Normative). The declared $\sigma(u,d,a)$ method MUST include the material elicitation, measurement, coding, model, and anchor-resolution uncertainty relevant to the cell claim. A zero uncertainty value is permitted only for an exact deterministic input whose transformation and reference are themselves exact for the declared claim. Where material resolution is unknown or unsupported, the run MUST not manufacture decisiveness from an arbitrarily small denominator; it must widen uncertainty, narrow the claim, or treat selection as non-decisive.

Method C restriction (Normative). Method C may be used only with a pre-registered calibrated table. A bare constant σ applied to all active cells is permitted only for worked examples, training artifacts, or explicitly non-operational demonstrations. A DEMONSTRATION_ONLY, UNVALIDATED, or materially sensitivity-flipping uncertainty posture MAY display point scores and a demonstration Gap, but MUST NOT support ALLOW_FRAMEWORK_SELECTION, SELECTED_DECISIVE, or an unqualified operational claim of unique framework selection.

Define approximate RLS uncertainty:

$$\sigma_{\text{RLS}}(a) = \sqrt{[\sum_u \sum_d (q(u,d) \times \sigma(u,d,a))^2] / Q}, \text{ for } Q > 0$$

This is the independent-cell approximation. Tier 3 and decision-material Tier 2 runs MUST declare dependence clusters and compute the perfect-correlation or covariance-aware stress specified by the Dependence-cluster uncertainty rule in Section 10.3. If ranking or decisiveness changes, set RLS_DEPENDENCE_SENSITIVE and treat the result as non-decisive until governed resolution. The declared cell uncertainty MUST also include material measurement or elicitation resolution; exact zero is allowed only for exact deterministic inputs.

Dependence-cluster uncertainty rule (Normative). The formula above is the independent-cell approximation. Because cells may share sources, causal pathways, stakeholder populations, scenario assumptions, models, or evaluators, the PCC MUST declare uncertainty-dependence clusters for Tier 3 and decision-material Tier 2 runs. At minimum, compute a perfect-correlation stress within each declared cluster by summing absolute weighted uncertainty contributions inside the cluster, then combine independent clusters by root-sum-square. Where covariance estimates exist, a covariance-aware calculation MAY replace the stress approximation. If ranking or decisiveness changes relative to the independent-cell result, set

RLS_DEPENDENCE_SENSITIVE and treat the result as non-decisive until the tie-break or governance review is complete.

Selection-claim materiality clarification (Normative). Any Tier 2 run that emits or publicly claims ALLOW_FRAMEWORK_SELECTION, SELECTED_DECISIVE, or an equivalent unique framework selection is decision-material for the dependence requirement above. The unique-selection claim must survive every required declared dependence treatment; the independent-cell result is an optimistic diagnostic rather than a substitute for cluster or covariance analysis. When material dependence cannot be credibly bounded, use the within-cluster perfect-correlation stress and refuse unique selection if the leader or decisive/non-decisive classification changes.

Optional risk-adjusted score:

$$\text{RLS_adj}(a) = \text{RLS}(a) - \lambda \times \sigma_{\text{RLS}}(a)$$

Default: $\lambda = 0.5$. Use is optional but must be declared.

By default, RLS_adj is diagnostic only and MUST NOT replace RLS for ranking or decisiveness unless the run declares ex ante that risk-adjusted ranking is the governing ranking basis. If such a declaration is made, the same declared basis MUST be used consistently for ranking, Gap, and decisiveness across all compared options in the run.

10.3A Dependence-Bound and Numerical-Guard Reporting (Normative)

Every Tier 3 and high-stakes Tier 2 decisiveness report SHALL identify the uncertainty aggregation method, dependence clusters, and the reason the method is appropriate. The independence or weak-dependence estimate MAY be reported, but it cannot be the sole decisiveness basis where shared evidence, common causes, common measurement instruments, or cross-cell pathways make dependence material.

Unless a covariance-aware model with declared provenance controls, report both the declared estimate and a conservative dependence bound or stress. SELECTED_DECISIVE requires $\text{Gap} > \delta$ under every required dependence treatment. A reversal sets RLS_DEPENDENCE_SENSITIVE.

The ϵ term in Gap is a numerical division guard, not empirical evidence and not hidden certainty. Reports SHALL state its value and whether it materially affects the denominator. If ϵ changes the decisive/non-decisive classification, set GAP_NUMERICAL_GUARD_SENSITIVE and refuse unique selection until the treatment is governed. Although numerical in purpose, a positive epsilon also regularizes the denominator on the declared score scale and can change a decisiveness classification. It is not evidence of uncertainty or certainty; retain the required numerical-guard sensitivity disclosure and do not alter epsilon after viewing the preferred result.

10.4 Discrimination Threshold (Decisive vs Non-Decisive) (Normative)

Stress-threshold reading rule. Let $V = \sigma_a^2 + \sigma_b^2$ under the declared independence-form arithmetic. A common uncertainty multiplier k requires $\text{DeltaScore} > \delta \sqrt{k^2 V + \epsilon}$. The corresponding base Gap requirement is $\delta \sqrt{(k^2 V + \epsilon)/(V + \epsilon)}$; it is approximately $k \delta$ only when epsilon is negligible relative to V . Neither δ nor this proxy Gap is a statistical z-score. All applicable dependence and robustness rules still apply.

Define gap between two options:

$$\text{Gap}(a,b) = \text{abs}(\text{RLS}(a) - \text{RLS}(b)) / \sqrt{\sigma_{\text{RLS}}(a)^2 + \sigma_{\text{RLS}}(b)^2 + \epsilon}$$

Default: $\delta = 2$, $\epsilon = 1e-6$.

Rationale (Informative). $\delta = 2$ is a conservative governance convention for the normalized Gap definition, not a claim that Gap is a literal z-score or that a specific distributional assumption holds. Its purpose is to reduce premature declarations of decisiveness when option-level uncertainty materially overlaps. Domains may tune

delta, but any non-default delta MUST be declared in the PCC, justified, and applied consistently across options in the run.

When at least two options are selectable, let a^* be the highest-scoring option under the declared governing ranking basis. $\text{Score}(a) = \text{RLS}(a)$ by default; $\text{Score}(a) = \text{RLS_adj}(a)$ only when permitted and declared ex ante under Section 10.3. For each other selectable option b , define $\text{SignedGap}(a^*, b) = [\text{Score}(a^*) - \text{Score}(b)] / \sqrt{\text{sigma_RLS}(a^*)^2 + \text{sigma_RLS}(b)^2 + \text{epsilon}}$. The pairwise absolute Gap remains a diagnostic; decisive selection uses the signed leader comparison.

Decisive lead rule:

- A unique decisive lead requires $\min_{\{b \in A_{\text{sel}}, b \neq a^*\}} \text{SignedGap}(a^*, b) > \text{delta}$ under every required comparison and uncertainty variant, the same leader in every required construction, and complete satisfaction of Section 10.4B. A missing, undefined or non-finite required comparison, or an incomplete contender/variant set, precludes unique selection; an empty comparison set cannot pass by default. These governed discrimination tests do not confer familywise statistical confidence. A sole surviving option is reported as a sole survivor without a fabricated pairwise Gap; qualification and authority requirements remain binding.
- If any selectable rival has $\text{SignedGap}(a^*, b) \leq \text{delta}$ in any required variant, or the leading option changes, the lead is non-decisive. Retain the unresolved contenders and apply Section 11.5; an intervening point-score runner-up MUST NOT conceal an unresolved lower-scoring rival. A permitted joint-uncertainty method must compare the leader against every rival over its declared joint set.

Cross-option dependence sensitivity (Normative clarification). The nominal Gap above is retained. Its root-sum-square denominator does not establish independence of option-estimation errors. Every Tier 3 or decision-material Tier 2 unique-selection claim SHALL assess whether common evidence, shared parameters or opposing allocation responses can create selection-material cross-option dependence. A computed material sensitivity at any tier remains subject to the cross-tier rule. Record the assessment and treatment in `GapSensitivityRecord` under the existing dependence robustness module; no new gate or decision state is created.

Where sigma values have a warranted final-score error interpretation and a supported joint-error model, a pairwise sensitivity may use $V_{\text{diff}}(a, b) = \text{sigma}_a^2 + \text{sigma}_b^2 - 2 \cdot \text{Cov}(\text{error}_a, \text{error}_b)$ and $\text{SignedGap_dep} = (\text{Score}(a) - \text{Score}(b)) / \sqrt{V_{\text{diff}} + \text{epsilon}}$. For positive sigma_a and sigma_b , $\text{Cov} = \rho_{ab} \cdot \text{sigma}_a \cdot \text{sigma}_b$ with ρ_{ab} in $[-1, 1]$. Use the adverse supported covariance bound, not an outcome-selected favorable correlation. Joint covariance constructions must be positive semidefinite and consistent with the declared marginals. If a sigma is exactly zero under its governing evidence rule, its covariance is zero; do not estimate a correlation by dividing by zero. A materially negative variance indicates invalid inputs/model, not a quantity to silently clamp into a decisive result.

For material but unbounded cross-option dependence, do not assume zero covariance. A conservative pairwise stress uses the lower correlation bound $\rho_{ab} = -1$, giving $V_{\text{diff}} = (\text{sigma}_a + \text{sigma}_b)^2$. This is a separate worst-case bound for each pair, not a claim that all pairs jointly have correlation -1. Non-statistical interval or heuristic proxies require a warranted joint interval/bound sensitivity on the final score scale; a covariance label alone cannot validate them. A missing warranted treatment blocks a unique-selection claim. Retain the nominal comparison and all required adverse variants; a reversal sets the existing `RLS_DEPENDENCE_SENSITIVE` and the result remains non-decisive under §§10.4A–10.4B. Favorable covariance cannot be used alone to rescue failed nominal decisiveness. The variance identity is standard uncertainty algebra, not evidence that RippleLogic estimates are empirically calibrated (NIST Technical Note 1297, Appendix A).

10.4A Non-Decisive Gap Enforcement (Normative)

If the leading options are tied, close, uncertainty-overlapped, or sensitivity-flipping under the declared discrimination rule, `FrameworkVerdict` MUST NOT state or imply deterministic framework selection. If any required $\text{SignedGap}(a^*, b) \leq \text{delta}$, or if any actually computed governed sensitivity changes the decisive/non-decisive classification, or if the uncertainty basis is limited to an unvalidated demonstration that cannot support the claimed operational strength, the run MUST either apply a governed tie-break rule over already-selectable

options, return REFUSE_DETERMINISTIC_SELECTION, or route to AuthoritySelectionRecord without claiming that RippleLogic itself selected uniquely.

Any run that claims deterministic framework selection despite one of those non-decisive conditions SHOULD set DECISIVENESS_OVERRIDE_INVALID with severity INVALID. Authority selection may still choose among admissible non-decisive options, but it must be recorded as authority selection, not framework determinism.

10.4B RLS Robustness Completeness (Normative)

Before ALLOW_FRAMEWORK_SELECTION, SELECTED_DECISIVE, or equivalent unique technical-selection language, the PCC SHALL contain one RLSRobustnessRecord drawn from the finite Canonical Robustness Module Registry.

Each module SHALL be in exactly one state: NOT_TRIGGERED_WITH_RATIONALE, PASS_NO_REVERSAL, SENSITIVE, UNRESOLVED, or REQUIRED_NOT_EVALUATED. NOT_TRIGGERED_WITH_RATIONALE SHALL identify the trigger examined, evidence basis, reviewer status, and reopen condition. High-stakes non-trigger determinations require independent review where practicable.

The registered module families are: effect-token and WDBIP completeness; rights non-dilution; partition and temporal-token integrity; likelihood semantics and probability order; adverse confidence; distributional/subgroup closure; reference family; baseline family; comparison-mask symmetry; interaction/non-separability; dependence and Gap guard; weight profile; option-set closure; MPS hypothesis; phantom/unknown; PLSS prominence; kernel/saturation; and any domain-mandated module.

Unique framework selection is available only when every required module is PASS_NO_REVERSAL or NOT_TRIGGERED_WITH_RATIONALE, no module is SENSITIVE, UNRESOLVED, or REQUIRED_NOT_EVALUATED, and the same leading option is decisively separated from every other selectable option under Section 10.4 in every completed selection-material variant.

Admissible-alternative bound. An alternative construction enters the robustness set only when it is evidence-supported, selection-material, independently motivated rather than invented after outcome inspection, and accepted through the declared adjudication process. Rejected alternatives SHALL be logged with rationale. This closes both under-testing and indefinite adversarial over-testing.

Joint-stress rule (Normative clarification). Required uncertainty-scale variants SHALL be evaluated within each required dependence treatment using that treatment's governing joint model or bound. Other required robustness perturbations SHALL also be evaluated jointly when they are mutually compatible, can plausibly coexist and their interaction can be selection-material. Do not double-count the same uncertainty source or combine logically incompatible models. Record the joint constructions and justified exclusions in the existing RLSRobustnessRecord and GapSensitivityRecord under the admissible-alternative bound; a material unresolved interaction blocks unique selection. This requires justified joint coverage, not an indiscriminate Cartesian product of all modules.

A sensitive or incomplete record does not erase the RLS calculation. It limits the claim to conditional ranking, non-decisive comparison, or authority selection. Later authority choice SHALL remain separate.

Module state	Meaning	Unique selection?
NOT_TRIGGERED_WITH_RATIONALE	Trigger examined; evidence, reviewer and reopen condition recorded	Yes, if justified
PASS_NO_REVERSAL	Required variant completed; leader and decisiveness stable	Yes
SENSITIVE	Leader or decisive status changes	No

Module state	Meaning	Unique selection?
UNRESOLVED	Material competing constructions remain	No
REQUIRED_NOT_EVALUATED	Module triggered but not completed	No

10.5 RLS Output Requirements (Normative)

PCC MUST record, per option in A_sel:

- $I_{prop_welfare}(u,d,a)$ (or a retrievable summarized representation)
- $RLS(a)$ and $\sigma_{RLS}(a)$ and the uncertainty method
- Ranking and whether decisive (Gap vs δ)

Human-facing numeric rendering (Normative presentation rule). Machine records and calculation traces SHALL retain the precision needed for deterministic replay. Ordinary human-facing summaries SHOULD render RLS, CVaR, normalized Gap, uncertainty, and derived scores to no more than three decimal places unless additional displayed precision is decision-material and its evidential basis is stated. Rounded displays MUST identify the controlling full-precision calculation record and MUST NOT be used to reconstruct a different gate or selection result.

SECTION 11: TIE-BREAKS, HOLLOWING DIAGNOSTICS, AND STRUCTURAL MONITORING

11.0 Purpose (Normative)

This section defines the special-case role of UCI/HOI after the core method has already formed a selectable set. UCI/HOI are not public headline cascade stages and are not independent late-stage admissibility gates. Their gate-relevant structural work belongs in Containment. Their residual work occurs only when RLS is tied, close, uncertainty-overlapped, or non-decisive, or when monitoring and hollowing-risk documentation are required.

- UCI is a structural metric used primarily as CSV evidence where material and secondarily as a residual tie-break signal when RLS cannot decide.
- HOI is a hollowing and drift diagnostic that flags welfare-up/coherence-down patterns for Containment, monitoring, and residual tie-break contexts.

11.0A Instrument-Placement Rule for UCI/HOI (Normative)

UCI and HOI are structural diagnostic instruments with two possible uses.

1. Containment use. When UCI/HOI evidence is material to cross-scale structural integrity, it is evaluated inside the CSV Gate under Section 9. This includes severe coherence loss, hidden hollowing, institutional brittleness, burnout, trust collapse, ecological degradation, governance-capacity loss, lock-in, or other structural damage to a containing union.
2. Residual tie-break use. When no containment threshold is crossed and options have already passed RF/NCRC, TRC, and CSV, UCI/HOI may be used after RLS only to resolve tied, close, uncertainty-overlapped, or non-decisive welfare rankings, document residual structural preference, trigger monitoring, or inform future-run steering.

Level 5 public-facing language SHALL NOT present UCI/HOI as a regular cascade stage. The public Level 5 is RippleLogic Score. The UCI/HOI rule is a special-case tie-break and monitoring rule attached to RLS non-decisiveness and CSV diagnostics.

11.1 Union Coherence Index (UCI) (Normative)

11.1.1 UCI components (canonical)

For each union u , define component scores in $[0,1]$:

Component	Symbol	Description
Cohesion	H_u	Internal connectivity, trust, shared identity, conflict resolution capacity
Flow	F_u	Coordination throughput, information fidelity, resource allocation efficiency
Resilience	R_u	Redundancy, robustness, recovery speed, adaptive capacity
Equity	E_u	Fair distribution of burdens/benefits, voice representation, inclusion

Let A_u be the set of components that are substantively applicable and measured for union scope u . Compute:

$$UCI_u = [\sum_{j \in A_u} \alpha_j * X_{\{u,j\}}] / [\sum_{j \in A_u} \alpha_j]$$

where $X_{\{u,j\}}$ is H_u , F_u , R_u , or E_u . Default weights are $\alpha_H = \alpha_F = \alpha_R = \alpha_E = 0.25$ when all four components are applicable. An inapplicable component is recorded as NA and the remaining weights are renormalized; it MUST NOT be assigned a perfect score merely to avoid a penalty.

For Self (U1), equity-as-distribution is ordinarily NA; the default computation therefore renormalizes across Cohesion, Flow, and Resilience unless a domain-specific within-self equity construct is explicitly defined and measured. For Biosphere (U7), Equity is likewise NA unless a defensible ecological-distribution or justice indicator is declared. Missing but applicable evidence is UNKNOWN, not NA, and triggers the applicable UCI-unavailability or claim-narrowing rule.

11.1.2 Structural independence rule (Normative; Tier 3 binding)

UCI MUST be computed from structural/process indicators distinct from welfare indicators used for RLS.

At Tier 3:

- Deriving UCI from welfare-cell impacts is PROHIBITED.
- If structural indicators are unavailable, UCI is treated as unavailable and escalation/judgment-call protocol must be used (Section 11.5; Appendix E).

11.1.3 ΔUCI computation (Normative)

11.1.3A UCI Measurement Maturity Statement (Normative acknowledgement)

Structural coherence and viability are load-bearing in the RippleLogic architecture. UCI is one governed diagnostic for those properties and a residual tie-break interface; it is not a universally mandatory measured input and does not by itself define CSV. A run MAY resolve CSV from direct, domain-appropriate structural evidence without computing UCI where UCI is not material, provided every applicable CSV subcondition, dependency, control, reversibility, monitoring, authority/legitimacy, carrier-capacity, and evidence-sufficiency requirement is explicitly resolved. Where UCI is material, it informs rather than substitutes for the structured CSV evidence case. The current release line specifies UCI component definitions (Cohesion, Flow, Resilience, Equity), structural independence requirements, and indicator families per scope, but does not yet provide universally validated measurement instruments with published psychometric properties (reliability, validity, normalisation benchmarks) across domains.

Scope	Measurement Maturity	Notes
U1 Self	Provisional	Validated instruments exist for some components (e.g. psychological resilience scales,

Scope	Measurement Maturity	Notes
		self-efficacy measures) but are not yet assembled into a UCI-specific battery.
U2 Household	Provisional	Relationship quality instruments exist; household-level aggregation rules are not yet standardised.
U3 Community	Provisional, strongest starter coverage	Social capital instruments (e.g. trust surveys), participation parity measures, and mutual-aid readiness metrics exist; normalisation and inter-rater reliability under UCI aggregation are untested.
U4 Organisation	Provisional, strongest starter coverage	Psychological safety instruments (e.g. Edmondson 1999), process throughput, incident recovery, and equity indices exist; assembly into UCI and structural independence verification are untested.
U5 Polity	Provisional	Institutional trust indices, rule-of-law indices, service delivery measures exist; UCI-specific aggregation is untested.
U6 Humanity / Global Coordination	Early design	Cross-polity cooperation and global coordination indicators are sparse and domain-specific.
U7 Biosphere	Early design	Ecosystem connectivity and resilience metrics exist (e.g. biodiversity integrity indices, trophic stability measures) but UCI aggregation is not yet specified. E7 is ordinarily NA and the applicable components are renormalized; a perfect score is never inserted by default.

Consequence for Tier 3 claims (Normative). Tier 3 run-level conformance claims are possible under the current specification when the run declares its UCI indicator sources, normalisation rules, and structural independence posture (Appendix E). Framework-level claims that UCI is empirically calibrated or cross-domain validated are PROHIBITED until a companion UCI Measurement Annex is published with at minimum: (i) named instruments per component per scope for U3, U4, and U5, (ii) normalisation rules with declared reference populations, (iii) structural independence verification evidence, and (iv) inter-rater reliability data (target ICC ≥ 0.70 on UCI component scores).

Validation commitment. Publishing such an annex is a Phase 2 validation priority (Section 17.3). Until it exists, the UCI measurement gap is the single largest barrier to framework-level empirical operational-readiness claims for Containment and tie-break governance.

Interpretive hardening (Normative). Until a companion UCI Measurement Annex is published, the structural questions UCI is intended to illuminate remain load-bearing, while UCI itself remains an empirically provisional diagnostic at the framework level. UCI therefore has no independent decision authority merely because a numerical value can be computed. CSV MAY resolve from direct, domain-appropriate structural evidence when UCI is not material, provided every applicable CSV subcondition is explicitly resolved; where UCI is material, it informs and challenges that evidence case rather than substituting for it. This limitation SHALL be surfaced in the Reader Guide, in Tier-claim language, and in any public claim about Containment or tie-break readiness. No framework-level claim of cross-domain calibrated or cross-domain validated UCI is permitted in the current v13.0 release line.

Provisional-UCI high-stakes restriction (Normative). A provisional UCI result MUST NOT be the sole evidential basis for either a high-stakes CSV_FAIL or an unqualified high-stakes CSV_PASS / CSV_PASS_WITH_CONTROLS, and MUST NOT by itself convert an otherwise non-decisive high-stakes RLS result into deterministic framework selection. A high-stakes fail requires a structured CSV evidence case that identifies the concrete structural mechanism, affected containing scope, evidence and uncertainty, threshold or failure condition, reviewer status, and available controls or redesign path. A high-stakes pass requires independently supported evidence that the relevant structural mechanism, dependencies, controls, reversibility, monitoring, and carrier capacity are adequate. If UCI is the only remaining discriminator among otherwise selectable options, the PCC SHALL disclose UCI measurement maturity, uncertainty, indicator provenance, structural-independence posture, and the underlying structural evidence; absent sufficient

support, preserve the non-decisive result and route any later choice through the separately identified authority/governance path. Where provisional UCI raises or clears a material concern but the structured evidence case is not yet sufficient, use ASSUMPTION_BOUND, CSV_ESCALATE, or CSV_REDESIGN_REQUIRED as appropriate rather than a falsely certain fail, pass, or deterministic selection.

For each option a:

$$\Delta UCI_u(a) = UCI_u(a) - UCI_u(\text{baseline})$$

Containment uses ΔUCI for containing unions (Section 9).

11.2 Aggregate UCI (Optional, informative)

An aggregate coherence score may be computed:

$$UCI_{agg} = \sum_u y_u \times UCI_u$$

Where y_u are declared aggregation weights. This aggregate is optional and MUST NOT replace per-union UCI in containment logic unless explicitly governed.

11.3 Hollowing-Out Index (HOI) (Normative as monitoring definition)

HOI detects a drift pattern: welfare score improves while coherence degrades over time.

Let t index review periods. Define:

- $\Delta RLS_t = RLS_t - RLS_{\{t-1\}}$
- $\Delta UCI_t = UCI_{agg,t} - UCI_{agg,\{t-1\}}$ (or a declared union-specific UCI; must be consistent)

Define exponential moving average EMA_λ with a declared smoothing parameter (default half-life 3 periods).

$$HOI_t = EMA_\lambda(\Delta RLS_t / s_{RLS})_t - EMA_\lambda(\Delta UCI_t / s_{UCI})_t$$

Scale rule (Normative). s_{RLS} and s_{UCI} MUST be positive, preregistered reference scales derived from a declared reference range or robust dispersion estimate and held fixed for the review window. If either scale is unavailable or unstable, HOI is UNAVAILABLE; raw unstandardized changes MUST NOT be subtracted.

Interpretation:

- Persistent $HOI > 0$ means standardized welfare change exceeds standardized coherence change; by itself it does not prove deterioration because both quantities may be improving. A strict hollowing pattern requires the declared EMA of standardized welfare change to be positive while the declared EMA of standardized coherence change is negative. When $HOI > 0$ but coherence is non-negative, report **RELATIVE_COHERENCE_LAG** rather than structural deterioration. HOI remains a monitoring diagnostic and never a gate or selection rule.
- HOI is diagnostic; it does not by itself make an option inadmissible, but it can trigger monitoring escalation and influence tie-break risk flags.

Normative limitation. HOI SHALL NOT by itself render an option rights-admissible, inadmissible, non-selectable, or rejected under Containment. Its normative use in the current release line is limited to monitoring escalation, reviewer caution, and tie-break risk signaling unless an explicit future version declares a stronger role.

11.4 Reserved for Future Structural-Diagnostic Extension (Normative numbering marker)

Section 11.4 is intentionally reserved so that established Section 11.5 and later cross-references remain stable. No operational rule is assigned to this marker in the current release line.

11.5 Canonical Tie-Break Chain (Normative)

This section applies only after the selectable set A_{sel} has been formed by Reality Grounding, RF/NCRC, TRC, and CSV, and only when the leading residual welfare options remain tied, close, uncertainty-overlapped,

or non-decisive under the Section 10.4 discrimination rule. The tie-break chain does not reopen NCRC, TRC, or CSV and is not part of the public headline method.

The governed tie-break order below applies only within the selectable set and under the Section 10.4 discrimination rule. If RLS is decisive, residual UCI/HOI is not required for selection, though it may still be recorded for monitoring or hollowing-risk documentation where appropriate.

Methods used here that can materially affect the outcome MUST be declared and applied consistently across compared options. Residual UCI or HOI SHALL NOT be used to compensate for unavailable containment evidence, failed Containment, insufficient Reality Grounding, or any prior gate failure. If structural evidence is gate-relevant, it belongs in Section 9 Containment, not in Section 11 tie-break.

Step 0: Ensure options under comparison are selectable (hold qualifying Reality Grounding, RF/NCRC, TRC and full CSV states).

If top RLS option fails containment, it is not in A_{sel} and must not be tie-broken into selection.

Step 1: UCI dominance rule.

Prefer the option with higher predicted UCI outcomes if the difference exceeds a governed UCI dominance threshold Δ_{UCI} .

Default: $\Delta_{UCI} = 0.05$ (must be recorded; may be tightened).

Operationally: compare the minimum-coherence-change among critical containing unions and/or the relevant union UCI profiles. The PCC must declare the tie-break UCI comparison method:

- Method UCI-A (default): Maximize minimum ΔUCI_u over containing unions relevant to the decision scope; or
- Method UCI-B: Maximize sum of ΔUCI_u over the directly affected unions; or
- Method UCI-C: Maximize declared aggregate UCI_{agg} change.

Whichever method is used MUST be declared and applied consistently across compared options.

UCI tie-break maturity restriction (Normative). At Tier 3 and in other high-stakes uses, a provisional UCI difference does not by itself authorize deterministic tie resolution. The PCC SHALL identify the structural evidence underlying the UCI comparison, its measurement maturity, uncertainty, indicator provenance, and independent-review status. If that evidence is insufficient, disputed, or the provisional UCI value is the only discriminator, retain REFUSE_DETERMINISTIC_SELECTION and proceed to Step 3 rather than manufacturing framework determinism from an unvalidated index.

Step 2: HOI risk flag (if monitoring context exists).

If one option is associated with persistent positive HOI in comparable deployments or in modeled trajectory, treat it as riskier and prefer the alternative if the risk is material, or escalate.

Step 3: Escalation / judgment call.

If UCI is unavailable or non-decisive and HOI does not resolve, the decision MUST escalate to:

- Additional data collection (structural indicators), and/or
- A higher tier, and/or
- A documented governance judgment call (PCC labeled JUDGMENT_CALL_TIEBREAK_NONDECISIVE) with explicit monitoring plan.

Serialization and authority boundary (Normative). A Step 1 or Step 2 preference is a governed tie-break, not RLS decisiveness. While the Section 10.4/10.4B unique-selection conditions remain unmet, FrameworkVerdict MUST remain REFUSE_DETERMINISTIC_SELECTION; record the declared tie-break method, preferred option, evidence and remaining uncertainty in the PCC. A tie-break alone MUST NOT emit ALLOW_FRAMEWORK_SELECTION or SELECTED_DECISIVE. A separately lawful AuthoritySelectionRecord may adopt that preference and serialize the decision as SELECTED_BY_AUTHORITY_NON_DECISIVE. Without that authority act, retain the recorded refusal or escalation disposition. Execution authorization remains a separate, current, scope-bound requirement. New

evidence may support a fresh decisive run only after the applicable qualification and robustness requirements are re-evaluated.

11.6 UCI Unavailability Rule (Tier 3) (Normative)

If UCI cannot be computed without violating structural independence:

- UCI MUST be treated as unavailable.
- If RLS lead is non-decisive, decision MUST either: (i) escalate for more structural data / higher tier, OR (ii) record a governance judgment call with explicit labeling and monitoring plan.

Audit label required: JUDGMENT_CALL_UCI_UNAVAILABLE.

11.7 Structural Safeguard Anti-Gaming (Normative)

- UCI MUST NOT be computed from RLS welfare impacts at Tier 3.
- UCI indicator choices must be documented; changes require versioning and are subject to challenge.

SECTION 12: MORAL STATUS, PROTECTION, PARTICIPATION, REALITY-MANAGEMENT CAPACITY, AND POWER: SGP v8.8 INTEGRATION (BINDING INTERFACE)

12.0 Purpose (Normative)

RippleLogic does not determine consciousness or moral patienthood on its own. It consumes only the governed interface outputs of the pinned Sentience Gradient Protocol v8.8. SGP outputs are evidence, profile, and status objects; they do not prove metaphysical consciousness or oneness, create legal status, reorder the cascade, establish species supremacy, or grant authority.

12.1 Type-Separation Rule (Normative)

The following objects MUST remain separate:

- 1) **MPS: Moral Patienthood Surface.** Evidence band, interval, confidence, and protection posture concerning possible welfare-bearing experience.
- 2) **FPP: Full Protection Plateau.** Constitutional or governed non-downgrade rights-of-protection status.
- 3) **GPR: Governance Participation Readiness.** The mode by which an entity may participate directly, with support, or through representation in a named process.
- 4) **SPR: Stewardship/Power Readiness.** Role-specific eligibility for a named consequential role under reality, rights, tail-risk, containment, competence, auditability, lawful-mandate, and revocability constraints.
- 5) **ICP: Intelligence Capacity Profile.** Informative evidence about learning, generalization, understanding, reasoning, planning, metacognition, robustness, and resource efficiency.
- 6) **RMCP: Reality-Management Capacity Profile.** Informative, multidimensional evidence about integrated reality-grounded modelling, recursive self-in-system representation, formal abstraction, cumulative knowledge, tool and intervention breadth, distributed coordination, long-horizon continuity, normative scope, causal reach, and responsibility tracking. **P100** is the open top-plateau status for functional equivalence to or greater integrated reality-management capacity than the current human collective calibration anchor.

Strict rule: MPS and FPP govern protection; GPR governs participation design; SPR informs role-specific authority review; ICP informs competence; RMCP informs integrated reality-management capacity and may support CMIU, GPR, or SPR analysis only through separately governed records. No one surface substitutes

for another. Low GPR, SPR, ICP, or RMCP MUST NOT reduce MPS or FPP. High ICP, RMCP, GPR, SPR, or P100 MUST NOT create sentience, FPP, CMIU membership, legal status, or authority. FPP MUST NOT create CMIU membership or authority.

12.2 SGP Authority and Version Pin (Normative)

RippleLogic Canon v13.0 pins SGP_v8.8.docx for its moral-patienthood, protection, participation, role-readiness, intelligence and reality-management-capacity interface. The local source hierarchy is in Appendix RELEASE. The prior repository source remains external and unchanged; remote publication is not claimed. v7.x SG_norm, SG_patient_norm, pillar scores and plateau/CMIU coupling remain legacy fields, not mathematical equivalents of the typed outputs; migration requires a documented current-protocol re-evaluation.

12.3 Human Full Protection Plateau (Normative; Non-Overridable)

For every human person H:

$FPP(H) := 1$

and every human welfare-impact instance uses:

$h_{\{k,welfare\}} := 1.0$

This is a constitutional anti-domination, equal-dignity, and full-inclusion rule, not an MPS evidence result or a claim that every human maximizes any consciousness or intelligence indicator. MPS(H,t) remains a separate evidence object and may be NOT_EVALUATED, MPS-NE, or otherwise evidence-governed without weakening $FPP(H)=1$ or $h_{\{k,welfare\}}=1$. The rule is independent of disability, age, illness, wakefulness, communication, productivity, or partial observability. No SGP or RippleLogic measurement may downgrade it.

Legacy compatibility note. Any historical field or display labelled MPSInterval_RL(H)=[1,1] SHALL be interpreted only as the fixed human welfare-inclusion set {1}. It is not an MPS evidence interval and MUST NOT be used as evidence that a human has been measured at a maximal consciousness level.

12.3A Humanity Collective RMCP-P100 Calibration (Normative Interface; Provisional Calibration)

Human FPP-100 and humanity RMCP-P100 are different claims. Every human person holds FPP-100. Separately, SGP v8.8 designates humanity collectively as the provisional current reference calibration anchor for integrated reality-management capacity because the human collective has demonstrated cumulative symbolic language, mathematics and science, technological intervention, intergenerational knowledge preservation, formal institutions, global coordination, planetary monitoring, Recursive Union Awareness, and causal reach across Earth systems.

The designation is a provisional, evidence-governed calibration construct, not a claim that every human individually has maximal capacity, that humanity governs wisely, that humans have greater intrinsic moral worth, or that P100 grants unrestricted authority. The open P100 plateau MAY also be assigned to biological, digital, hybrid, collective, or extraterrestrial intelligences that reach, exceed, or may already exist at functionally equivalent capacity, subject to a valid SGP v8.8 RMCP record and independent governed review. An entity that exceeds the human anchor remains in P100; the difference is represented in the profile rather than converted into a 101 supremacy caste.

12.4 Where SGP Enters RippleLogic (Normative)

Permitted interfaces:

1. **Protected-stakeholder and Rights Floor interface.** MPS evidence and protection posture may identify a possible protected stakeholder and trigger precautionary RF/NCRC review. FPP establishes a non-

downgradable protection floor. Absence of a pre-existing rights-coverage mapping is not permission for harmful treatment.

2. **Residual Welfare-stream interface.** A valid SGP v8.8 record supplies a typed MPS evidence band/interval, protection posture, and record validity status. RippleLogic uses that record to construct a governed welfare-inclusion hypothesis set H_E for robust RLS sensitivity; it **MUST NOT** treat the MPS interval itself as a cardinal welfare multiplier. The Base stream fixes $h_k=1$. Every human person and FPP entity uses $H_E=\{1\}$. For a non-FPP entity with realistic or credible welfare-bearing possibility, H_E **MUST** include 1; where exclusion remains evidentially admissible it also includes 0. MPS-NE requires $\{0,1\}$ when decision-material. Additional bounded coefficients require a separate governed and validated welfare-attention profile and cannot be inferred from the MPS band.
3. **Participation interface.** GPR may inform direct, supported, or represented participation in a named governance process. GPR is not moral worth, voting weight, office eligibility, or command authority.
4. **Power interface.** SPR may be admitted as one role-specific evidence object inside an AuthoritySelectionRecord, but it never replaces RG, RF/NCRC, TRC, CSV, domain competence, lawful mandate, institutional responsibility, or revocability.
5. **Competence interface.** ICP may inform domain-competence evidence for GPR or SPR. It cannot alter MPS or FPP by itself.
6. **Reality-management capacity interface.** A valid RMCP record may inform CMIU analysis, collective-capacity comparison, GPR design, or SPR role review. It **MUST** identify the evaluated individual, class, institution, collective, or distributed-system boundary; the declared domain and scale; each RMCP dimension and uncertainty interval; external scaffolding; causal reach; evidence dependencies; capacity-performance separation; P100 status if any; and revision triggers. RMCP, RMCi summaries, Recursive Union Awareness, and P100 **MUST NOT** alter MPS or FPP, act as cardinal moral-worth weights, or bypass lawful mandate and the cascade.

12.5 Interval-Sensitivity Rule (Normative)

If the preferred option, decisive/non-decisive status, subgroup result, or material welfare conclusion changes across the governed welfare-inclusion hypothesis set, the run **MUST** set `MPS_HYPOTHESIS_SENSITIVE`. The PCC **MUST** report the affected stakeholders, MPS evidence record, hypotheses tested, changed outputs, and disposition. Permitted dispositions are: collect additional evidence; adopt a precautionary policy acceptable across the hypothesis set; narrow the claim; preserve multiple options; or refuse a definitive ranking. Selecting a convenient intermediate coefficient or suppressing hypothesis sensitivity is non-conformant.

12.6 Prohibited Inferences (Normative)

The following are prohibited:

Claim	Status
Linguistic fluency or an AI assertion of consciousness, suffering, personhood, or protection entitlement proves sentience or moral patienthood	NOT VALID
Intelligence, understanding, compliance, or tool performance proves moral patienthood	NOT VALID
Reward, loss, utility, or refusal signals alone prove valence	NOT VALID

Claim	Status
Low agency, language, metacognition, GPR, or SPR reduces credible welfare protection	PROHIBITED
MPS is a probability of consciousness, quantity of experience, or intrinsic-worth score	NOT VALID
FPP implies GPR, CMIU membership, office, command, or stewardship authority	PROHIBITED
GPR implies SPR or lawful authority	NOT VALID
RMCP, Recursive Union Awareness, or P100 proves sentience, FPP, moral superiority, metaphysical oneness, or wisdom	NOT VALID
P100 is a permanent human monopoly or a licence to dominate beings below the plateau	PROHIBITED
RMCP or P100 creates CMIU membership, legal status, command, or stewardship authority	PROHIBITED
SPR can bypass RG, RF/NCRC, TRC, CSV, or lawful mandate	PROHIBITED
Missing, NE, incomplete, or inadmissible evidence means zero sentience or zero reality-management capacity	PROHIBITED
A v7 scalar may be silently translated into a v8.8 output	PROHIBITED
An AI denial of consciousness, sentience, welfare, or personhood, or a statement that it is "only a tool," proves phenomenal absence or absence of moral patienthood	NOT VALID
Instrumental tool-role proves low capability, absence of agency, or absence of possible sentience	NOT VALID

12.7 CMIU Governance Definition and Entry Boundary (Normative)

The Collective Managing Intelligence Union is the U6 governance-coordination view of intelligences, representatives, and institutions with a declared participation channel or lawful role in shared governance, coordination, and systemic-risk management. It is not a moral-status class and is not administered by SGP.

CMIU participation or membership, where an implementation uses that concept, requires an explicit institutional rule. For an entity, direct or supported participation SHOULD reference GPR; any consequential role MUST reference a role-specific SPR record and AuthoritySelectionRecord. A valid RMCP record MAY inform whether a candidate can model, coordinate, preserve, or manage shared reality at the relevant scale, but RMCP or P100 never supplies lawful mandate or membership by itself. Institutions may participate through lawful mandate and accountable representatives. MPS, FPP, RMCP, or P100 alone MUST NOT create CMIU membership.

12.8 Rights Expansion for Non-Human Stakeholders (Normative)

When SGP v8.8 supports a stronger protection posture or FPP for a non-human individual or class, RippleLogic may expand rights coverage only through a governed update: update C_r and subgroup protocols; map the typed SGP protection bundle to the applicable protected domains LIFE, BODY, NEED, DIGN, LBTY, PROC, INFO, and ECOL; version the decision; record membership and identity conditions;

preserve prior PCCs; and provide notice, review, representation, and appeal. Class priors may raise protection but MUST NOT downgrade atypical individuals or create FPP without governed assignment.

SECTION 13: WEIGHT GOVERNANCE: HYBRID DEMOCRATIC WEIGHTING (HDW)

13.0 Purpose (Normative)

RLS requires union and dimension weights. Weights encode normative prioritization among admissible options. Without governance, weights become a capture surface. HDW provides legitimate, anti-capture weight governance by combining:

- Constitutional floors (non-negotiable minimum protection attention)
- Democratic tuning (stakeholder voice)
- Structural evidence input (system-level constraints)

HDW affects only ranking among selectable options. HDW MUST NOT alter admissibility gates (NCRC, TRC) or CSV gate.

13.1 Constitutional Floors (Normative; prior floor vector preserved through RippleLogic v13.0)

13.1.1 Union weight floors

Status of these numbers (Normative). The floor values below are chosen, challengeable governance priors - not discovered empirical constants. Any revision requires a declared justification, a versioned change, and sensitivity testing against retained alternatives. Floors MUST NOT be tuned after observing an outcome in order to manufacture a preferred ranking.

Union	Floor
Self (U_1)	0.20
Household (U_2)	0.06
Community (U_3)	0.06
Organization (U_4)	0.06
Polity (U_5)	0.08
Humanity / Global Coordination (U_6)	0.10
Biosphere (U_7)	0.10
Total	0.66

13.1.2 Dimension weight floors

Dimension	Floor
Material (D_1)	0.08
Health (D_2)	0.10
Social (D_3)	0.08

Dimension	Floor
Knowledge (D ₄)	0.08
Agency (D ₅)	0.10
Meaning (D ₆)	0.06
Environment (D ₇)	0.10
Total	0.60

13.1.3 Floor meaning (Normative)

Constitutional weight floors are minimum attention constraints in residual welfare ranking among options that have already passed the relevant admissibility gates. They are not rights. Rights are handled through NCRC in Section 7. They are not predictions of welfare salience for any particular decision. They are anti-erasure governance priors.

A floor value means that the corresponding Union Scope or welfare dimension may not be reduced to zero attention in residual welfare scoring merely because it is politically weak, distant, non-vocal, non-human, future-facing, or inconvenient to model.

Floor weights do not make all scopes equally active in every decision. PLSS may concentrate residual welfare attention on prominent scopes where the local ripple is strongest, provided that NCRC, TRC, CSV, non-maskable cells, and constitutional floors remain preserved.

Floor weights do not override evidence. They preserve representation. They ensure that optimization cannot erase a scope or dimension before the system has accounted for whether rights, tail-risk, containment, or welfare impacts are present.

Design logic. The canonical floor values are governance priors chosen to prevent erasure while leaving residual weight available for context-sensitive allocation. They are subject to governed revision only through the versioned update process, sensitivity analysis, and release-manifest discipline. They SHALL NOT be modified inside a run after outcome evidence is observed.

13.1.4 Floor derivation rationale (Normative governance prior; Informative derivation)

The constitutional floor values are governance priors designed to satisfy three structural requirements simultaneously: anti-erasure (no scope or dimension may be reduced to zero attention in welfare ranking among admissible options), proportionality (scopes with longer characteristic timescales, broader stakeholder coverage, or stronger enabling-condition roles receive moderately higher floors), and residual allocation headroom (floor totals are kept well below 1.0 to preserve meaningful democratic and structural tuning space).

Specific design logic for union floors:

Self (U1) = 0.20. Highest floor because individual-level experience, health, agency, and concentrated burdens are the constitutionally designated base welfare surface from which higher-scope aggregates are constructed. The floor prevents person-level impacts from disappearing inside household, organizational, polity, humanity, or biosphere summaries. It is not a decision-maker self-interest prior: a decision-maker or operator enters U1 only when materially affected and mapped under the same stakeholder and redundancy rules as anyone else. The floor remains below 0.25 so broader scopes collectively retain majority residual allocation.

Household (U2) through Organisation (U4) = 0.06 each. Set as equal modest floors reflecting that these scopes vary substantially in relevance across decision classes, so residual allocation rather than floor values should drive their relative emphasis.

Polity (U5) = 0.08. Modestly elevated above U2 through U4 because procedural legitimacy, enforcement capacity, and rule-of-law infrastructure are enabling conditions for all other scopes.

Humanity / Global Coordination (U6) and Biosphere (U7) = 0.10 each. Elevated because these scopes carry the longest characteristic timescales, the broadest stakeholder coverage, and the strongest enabling-condition roles. Degradation at these scopes is often irreversible and threatens the viability of all contained scopes.

Specific design logic for dimension floors:

Health (D2), Agency (D5), and Environment (D7) = 0.10 each. Elevated because these dimensions are most directly tied to rights-floor and enabling-condition protection: Health to survival and functioning, Agency to autonomy and non-domination, Environment to ecological viability.

Material (D1), Social (D3), and Knowledge (D4) = 0.08 each. Set as equal moderate floors reflecting their cross-cutting importance without prioritising any one.

Meaning (D6) = 0.06. Lowest floor because Meaning is the dimension most resistant to reliable cross-cultural measurement (Section 19.1) and most context-dependent in its governance relevance. The floor is non-zero to prevent erasure but is set conservatively pending measurement improvement.

Status (Normative). These floors are governance priors, not empirically derived constants. They are explicitly challengeable under the same governance rules that apply to all RippleLogic normative parameters: revision requires justification, version increment, PCC documentation preserving prior values, and sensitivity outcomes over a reference decision suite (Section 7.1). The values have remained stable since v6.0 and are carried forward in SGP v8.8 and preserved through RippleLogic v12.0 because no revision proposal has been submitted that satisfies these requirements.

Sensitivity guidance (Normative). Tier 3 deployments SHOULD include floor perturbation in their sensitivity bundle, testing at minimum ± 0.02 on each floor value (clamped to maintain $\text{sum} \leq 1$ and all floors ≥ 0.02). If selection changes under floor perturbation, the PCC MUST document this sensitivity.

Floors prevent unions and dimensions from being reduced to zero by weight governance. Active-cell rules, applicability masks, and declared scope structure separately determine whether a given scope or dimension contributes to a specific run. Floors are not rights, because rights are handled by NCRC; floors are minimum attention constraints in welfare optimization among admissible options. The floor values specified below are the canonical v6.0 floors, carried forward in SGP v8.8 and preserved through RippleLogic v12.0, and preserved unchanged in subsequent releases unless explicitly amended in the version history.

13.2 HDW Blend Formula (Normative)

Let:

- w^{floor} be union floors,
- v^{floor} be dimension floors.

Define allocable mass:

- $\text{allocable_U} = 1 - \sum_u w^{\text{floor_u}} = 0.34$
- $\text{allocable_D} = 1 - \sum_d v^{\text{floor_d}} = 0.40$

Let w^{dem} and w^{str} be union proposal vectors on the simplex ($\sum_u w^{\text{dem_u}} = 1$, all nonnegative). Let λ_U be the democratic share of allocable union mass. Default: $\lambda_U = 0.70$.

Clarification (Normative). Proposal vectors are constrained to the probability simplex: all components nonnegative, sum to 1, and subject to any declared floor constraints before blending. Validators MUST reject proposal vectors that violate simplex or floor constraints.

Then:

$$w_u = w^{\text{floor_u}} + \text{allocable_U} \times (\lambda_U \times w^{\text{dem_u}} + (1 - \lambda_U) \times w^{\text{str_u}})$$

Similarly for dimensions with λ_D (default $\lambda_D = 0.70$):

$$v_d = v^{\text{floor}_d} + \text{allocable}_D \times (\lambda_D \times v^{\text{dem}_d} + (1 - \lambda_D) \times v^{\text{str}_d})$$

By construction:

- $\sum_u w_u = 1$,
- $\sum_d v_d = 1$,
- All floors are satisfied,
- Weights remain nonnegative.

13.3 Democratic Proposal Process (Normative minimum)

HDW requires a documented process for producing w^{dem} and v^{dem} . At minimum:

- Participants: Representatives from affected unions (including vulnerable populations),
- Method: Transparent vote or deliberative process with published results,
- Publication: Results recorded in PCC and (Tier 3 recommended) in an immutable ledger.

13.3A Democratic Tuning Procedure (Normative minimum)

At minimum, HDW democratic tuning SHALL use a declared participation protocol that produces reproducible proposal vectors for scopes and dimensions. The default ballot method is a two-ballot allocation: each eligible participant allocates 100 points across the seven operational scopes to produce a raw scope proposal, and 100 points across the seven welfare dimensions to produce a raw dimension proposal.

Normalization rule (Normative). Raw ballot totals SHALL be normalized onto the simplex to produce w^{dem} and v^{dem} . Ballots with missing, negative, or over-totaled entries MUST be corrected or excluded under a declared correction rule recorded in the PCC. The PCC MUST record participant eligibility, quorum rule, number of ballots counted, exclusion/correction rule, and the final normalized proposal vectors.

Minimum procedural safeguards (Normative). Tier 3 HDW runs MUST include: (i) declared eligibility criteria, (ii) quorum threshold, (iii) publication of aggregate results, (iv) challenger opportunity before final locking, and (v) retention of the raw ballot summary sufficient for independent recomputation. If these conditions are not met, the run MUST use a declared interim weight profile and MUST NOT present the democratic proposal as fully satisfied HDW governance.

Floor lock (Normative). Democratic tuning operates only over allocable mass above constitutional floors. No democratic proposal may reduce any scope or dimension below its Section 13.1 floor. For Tier 3 HDW runs, where a ballot-derived proposal would reduce a final combined weight w_u or v_d to at or within 0.02 above its constitutional floor, the Section 13.5 supermajority safeguard applies. Evaluate the final proposed weight after the declared HDW combination, not an unblended ballot share. Stricter applicable governance requirements remain binding.

13.4 Structural Proposal Process (Normative minimum)

w^{str} and v^{str} represent evidence-informed constraints (for example, externality reach, irreversibility, systemic risk). At minimum:

- Method: Documented rule for producing w^{str} , v^{str} ,
- Inputs: Declared indicators and sources,
- Reproducibility: The same method yields same output given same data.

If a deployment lacks structural evidence, it may use a declared interim default w^{str} , v^{str} at Tier 2, but Tier 3 SHOULD converge to a governed structural method over time.

13.5 Anti-Capture Safeguards (Normative for Tier 3)

Tier 3 weight governance MUST include:

- Stratified representation: Delegates must include vulnerable population representation and biosphere stewardship for environment-relevant decisions.

- Supermajority lock near floors: Any Tier 3 proposal reducing a final combined weight w_u or v_d to at or within 0.02 above its floor requires $\geq 2/3$ approval. The test uses the final proposed weight, not an unblended proposal share.
- Transparency ledger: All proposals, votes, rationales published (at least internally).
- Red-team testing: Proposed weights tested against reference decision suites to detect systematic bias.
- Conflict-of-interest disclosure: Material conflicts require recusal; recusals recorded.

13.6 Weight Use by Tier (Normative)

Tier	Weight Policy
Tier 1	Floor-preserving uniform residual allocation allowed (documented); raw 1/7 union weights are prohibited where constitutional floors apply.
Tier 2	Floor-preserving uniform residual allocation allowed; HDW recommended; weights must be declared in PCC.
Tier 3	HDW recommended; if not available, explicit interim weights allowed with justification and sensitivity analysis. Floors remain binding regardless.

13.6A Weight Constitution, Induced Cell Floors, and Normative Underdetermination (Normative)

Every claim-bearing RLS run SHALL declare WeightProfileStatus as AUTHORIZED, INTERIM, CONTESTED, or UNSETTLED. AUTHORIZED means a legitimate ex-ante process selected the controlling profile; the resulting ranking remains conditional on that constitution.

INTERIM requires sensitivity across the live legitimate profiles. CONTESTED or UNSETTLED requires comparison across those profiles when selection-material. Arbitrary post-result profiles are inadmissible.

Constitutional union and dimension floors are minimum anti-erasure attention constraints. Because $q(u,d)$ includes the product $w_u \cdot v_d$, they also induce minimum cell-level ranking influence. The induced cell floor is multiplicative, not additive, and reports SHALL disclose the resulting cell-floor range when it is material. This influence is a chosen constitutional convention, not an empirically discovered marginal welfare exchange rate.

If two or more legitimate ex-ante weight constitutions select different options after the empirical and methodological surface is sufficiently closed, return REFUSE_DETERMINISTIC_SELECTION with NonDecisiveReason=NORMATIVE_UNDERDETERMINED. Additional empirical evidence cannot determine which moral or political constitution is uniquely correct. An accountable authority MAY select among the disclosed selectable options through a separate AuthoritySelectionRecord.

SECTION 14: AUDITABILITY AND DETERMINISM: PCC AND AIL (TIER 1-3; TIER 4 TARGET)

14.0 Purpose (Normative)

Without auditability, ethics frameworks become theater or are gamed. RippleLogic therefore requires a structured decision artifact (PCC) and integrity rules (AIL) to make decisions reconstructable, challengeable, and corrigible via NCAR.

Tier boundary: Tier 1-3 auditability is fully specified here. Tier 4 determinism and hash-bound replayability are explicitly a design target only (Appendix I) and MUST NOT be claimed in the current v13.0 release.

14.1 Artifact Integrity Law (AIL) Principles (Normative for Tier 2-3)

AIL is a set of integrity constraints governing how decisions are recorded and compared.

Principle	Description
AIL1 (Registry Binding / Source Traceability)	Every Tier 2-3 PCC MUST list the normative parameters used. If hashes/registries are used, they must be referenced; if not, values must be embedded explicitly.
AIL2 (Immutability)	A PCC is immutable after signing. Any correction produces a new PCC revision that references the prior PCC and explains the change.
AIL3 (Comparability)	Options compared in one run MUST be evaluated under identical configuration.
AIL4 (No Silent Overrides)	Any override must be explicit in the PCC (what changed, why, who approved).
AIL5 (Auditability Sufficiency)	A Tier 3 PCC MUST contain enough information for an independent reviewer to recompute NCRC, TRC, CSV, and RLS for that run.

14.2 PCC Requirements by Tier (Normative)

PCC.EKRegistry (CONDITIONAL; REQUIRED when any non-default e_k is used, or when a Tier 2 named-instrument exception is invoked). For each instrument record: InstrumentID; Purpose; Source/Authority; AllowedRange; TierPermitted; ReasonCodeSet; DefaultCounterfactualRequired [TRUE/FALSE].

Tier	PCC Requirement
Tier 1	Optional but recommended for learning
Tier 2	REQUIRED (basic) including Scope Coverage Declaration and Stakeholder Coverage Index ($SCI \geq 1$). A full Stakeholder Discovery Protocol (SDP) record is REQUIRED when any SDP trigger holds (Section 2.2A); otherwise the PCC MUST include either a lightweight stakeholder note stating that no trigger held and why, or an SDP record voluntarily. InstanceMap is REQUIRED when any SDP trigger holds (Section 2.2A).
Tier 3	REQUIRED (full) including subgroup rights checks, TRC scenario table, containment results, sensitivity analysis bundle, Scope Coverage Declaration, Stakeholder Coverage Index ($SCI \geq 2$), and full InstanceMap with multi-scope redundancy handling disclosures.

14.3 Audit Flags (Normative)

Audit flags are standardized labels for integrity failures or risk warnings. Flags MUST be recorded in PCC when triggered. If any INVALID flag triggers, the PCC is invalid and the decision run MUST be recomputed after correction.

Definitions for this section:

- Mode A (Containment as binding gate): Containment / CSV evaluation determines which options are selectable; options failing CSV MUST NOT be selected. See Section 9 and Appendix B.11.
- Mode B (diagnostic-only legacy posture): a structural diagnostic is computed for information but does not affect selection. Mode B is PROHIBITED for any ordinary ranking or selection claim. Current selection-relevant structural evidence belongs inside CSV under Section 9; residual non-gate-critical use is limited by Sections 10.4 and 11.
- KQS (Kernel Quality Score): Summary score indicating kernel readiness for propagation. See Section 6.5.
- Kernel perturbation test: Systematic edge perturbation of ± 0.05 (or $\pm 10\%$ of magnitude, whichever is larger) applied one-at-a-time to relied-upon non-zero kernel edges. See Section 6.6.

Canonical Audit Flags - Current v13.0 Specification

14.3A Canonical Audit Flag Registry v13.0 (Normative)

Section 14.3A and docs/implementation/CANONICAL_AUDIT_FLAG_REGISTRY_v1.1.yaml are one controlling audit-flag registry expressed in human-readable and machine-readable form. The YAML controls token enumeration, owner, severity vocabulary, and source location; this section controls normative meaning and required action. A controlling token introduced anywhere in the Canon or an active companion MUST be

added to the YAML and projected here or be explicitly labeled local/informative and non-controlling. The release verifier MUST fail when a normatively controlling token is absent from the machine registry or when the two surfaces assign incompatible ownership or severity.

Reading-package and machine-release boundary (Normative clarification). Every controlling token explicitly required by this Canon retains its meaning and required disposition even when an external registry is unavailable. Absence from a mirror or unavailable YAML MUST NOT be treated as permission to suppress it. Publication claiming registry-conformant machine interchange remains blocked until the complete YAML projection and its source meanings are manifest-pinned, reconciled and verified. A bounded research/reading package that explicitly disclaims that stronger conformance does not establish the missing external registry by assertion. Package-local derived directories are navigation/testing aids, not silently promoted replacements for that authority.

14.3A.1 Supplemental current-line ownership registry (Normative)

Ownership rule. Canon tokens govern framework conformance. Agent-runtime tokens are implementation events mapped to Canon states; they do not create new gates or modify Canon precedence.

RUNTIME_BLOCK and RUNTIME_INVALID are runtime-action severities owned by the Agent System; they map respectively to execution blocking/escalation and invalid runtime-state claims. They do not extend the Canon PCC severity vocabulary {INVALID, ESCALATE, REVIEW, DISCLOSE}.

Token	Trigger	Required action	Severity	Owner
DECISIVENESS_OVERRIDE_INVALID	A non-decisive result is overridden without the required lawful authority and separate AuthoritySelectionRecord.	Invalidate the deterministic-selection claim; restore REFUSE_DETERMINISTIC_SELECTION and record any lawful authority selection separately.	INVALID	Canon
FALSIFICATION_LAYER_UNATTRIBUTE D	A falsification, revision, or failure claim is not assigned to the affected dependency layer.	Attribute the failure layer and reopen all materially dependent surfaces.	ESCALATE	Canon
FALSIFICATION_LAYER_SKIPPED	A materially affected downstream layer is not reopened after an upstream failure.	Reopen and rerun the skipped layer or narrow/withdraw the dependent claim.	INVALID if claim-bearing	Canon
POST_HOC_PROJECTION_INVALID	Projection, mapping, threshold, or evaluation rule is changed after results without a versioned revision branch.	Invalidate the affected comparison and rerun from the frozen evaluation boundary.	INVALID	Canon
NON_UNIQUE_PREDICTION	A claimed necessary prediction is not unique under the declared premises or alternatives.	Narrow to an assumption-bound or exploratory claim and record alternatives.	REVIEW / ESCALATE	Canon
EXPLORATORY_ONLY	Evidence or method supports exploration but not a confirmatory or decision-authorizing claim.	Label exploratory; prohibit confirmatory, selection, authority, or deployment wording.	DISCLOSE	Canon

Token	Trigger	Required action	Severity	Owner
OUTPUT_WITHOUT_ADMISSIBLE_REFERENCE	A consequence-bearing output is requested or produced without an admissible reference.	Refuse, narrow, or escalate; do not execute.	RUNTIME_BLOCK	Agent runtime
STRUCTURAL_VIABILITY_FAIL	The runtime StructuralViabilityRecord resolves fail.	Block execution and route to redesign or escalation.	RUNTIME_BLOCK	Agent runtime
SVC_INPUT_INSUFFICIENT	Required structural-viability input is absent or unresolved.	Remain advisory-only; request evidence or escalate.	RUNTIME_BLOCK	Agent runtime
PERMISSION_WITHOUT_ADMISSIBILITY_INVALID	Permission or approval is used as a substitute for admissibility.	Block execution and record the category collapse.	RUNTIME_INVALID	Agent runtime
STRUCTURAL_OVERRIDE_OF_RIGHTS_INVALID	A structural or operational consideration is used to override RF/NCRC.	Block execution; RF/NCRC controls.	RUNTIME_INVALID	Agent runtime

Flag	Trigger	Required action	Severity
CATEGORY_GROUNDING_RECORD_MISSING	Category Grounding is required and PCC.CategoryGroundingRecord is missing.	Block deterministic framework selection, conformance, public-deployment, and optimized-metric claims until the claim is narrowed, escalated, marked sensitivity-only / monitor-only, or refused.	INVALID
SEMANTIC_DEBT_AGGREGATE_HIGH	Semantic debt score reaches or exceeds the governed threshold for the claim context.	Require reviewer disposition before stronger public, conformance, authority-selection, deployment-readiness, or deterministic-selection claims.	ESCALATE
SILENT_REMAINDER_INVALID	A required field is missing or an undefined remainder is silently treated as pass, zero, or non-material.	Declare the remainder, collect evidence, rerun, narrow the claim, escalate, or refuse.	INVALID
INSUFFICIENT_GROUNDING	The declared reality surface is insufficient for a material rights, tail-risk, CSV, authority, conformance, or selection claim.	Narrow the claim, collect evidence, rerun, mark exploratory/sensitivity-only, escalate, or refuse the stronger	ESCALATE/REFUSE

Flag	Trigger	Required action	Severity
		claim.	
CONTAINMENT_STRUCTURAL_DIAGNOSTIC_UNAVAILABLE	Structural diagnostic evidence required for CSV, including UCI/HOI or declared equivalent where material, is unavailable or unresolved.	Downgrade claim strength, collect evidence, rerun, escalate, refuse the stronger claim, or record permitted waiver where allowed.	ESCALATE
STRUCTURAL_DIAGNOSTIC_MISPLACED_TO_TIEBREAK	A structural diagnostic that was gate-relevant to CSV was treated only as residual tie-break evidence after RLS.	Rerun CSV with the diagnostic included or refuse the stronger claim.	INVALID for Tier 3; ESCALATE otherwise
PLSS_STRUCTURAL_DIAGNOSTIC_OMISSION	A PLSS run omitted, down-weighted, or deferred structural diagnostic evidence that was plausibly CSV-relevant.	Rerun with broader scrutiny, document non-relevance, or escalate.	ESCALATE
GATE_CRITICAL_CONFIDENCE_UNDERBOUNDED	A low-confidence adverse impact in C _r , C _{cat} , or a material CSV indicator is used to reduce apparent severity such that RF/NCRC, TRC, or CSV would pass when a conservative bound, rerun, downgrade, escalation, or Refusal is required.	Apply the gate-critical confidence guard; record conservative bound or unresolved gate disposition in PCC.	ESCALATE
SGP_BINDING_INTEGRITY_FAIL	A run claims SGP-aware protection, welfare-inclusion hypothesis sensitivity, FPP, GPR, SPR, or ICP use but omits the valid version-pinned record, target boundary, band/interval, hypotheses evaluated, or required registry fields; or it silently falls back to MPS-0, zero, a convenient coefficient, or a legacy v7 scalar.	Mark the affected PCC claim INVALID; restore the valid record and governed hypothesis set, rerun sensitivity, or remove/narrow the SGP-aware claim. No silent fallback is permitted.	INVALID for affected SGP-aware claim
RLS_DEPENDENCE_SENSITIVE	Ranking or decisiveness changes when uncertainty dependence clusters, perfect-correlation stress, or covariance-aware aggregation is applied.	Treat the result as non-decisive; disclose the dependence structure and complete the tie-break or governance review.	ESCALATE
SOURCE_COUPLING_RECORD_MISSING	Source-Coupling Integrity is triggered but the required SourceCouplingRecord is absent.	Add the SourceCouplingRecord, narrow the claim, downgrade to exploratory/monitor-only, escalate, or refuse the stronger claim.	INVALID for affected strong claim

Flag	Trigger	Required action	Severity
SOURCE_COUPLING_STATUS_UNSUPPORTED	Declared Source-Coupling status is stronger than the evidence, capability surface, or warrant supports.	Downgrade the status, add evidence, rerun RG/CSV if material, or invalidate the affected claim.	ESCALATE/ INVALID
PC_AEP_RECORD_MISSING	PC-AEP is triggered for a physical, causal, cyber-physical, medical, ecological, infrastructure, weapons, nuclear-command, planetary-cognition, or other high-consequence execution claim, but the profile is absent.	Add the PC-AEP, narrow the claim, mark governance-permission-only, escalate, or refuse the physical-execution claim.	INVALID for physical-execution claim
PC_AEP_STATUS_UNSUPPORTED	Declared Physical/Causal Admissibility status exceeds what the supplied domain evidence, validity domain, boundary conditions, verification surface, monitoring, shutoff, or expert warrant supports.	Downgrade status, collect domain evidence, narrow validity domain, rerun RG/CSV/TRC as needed, or refuse.	ESCALATE/ INVALID
PHYSICAL_EXECUTION_CLAIM_OVERSTATED	Governance permission, compliance, simulation success, fluent model output, approval, or an RLS score is represented as physical admissibility or physical safety proof.	Correct the claim language, require PC-AEP support, downgrade to governance-permission-only or not-established, and invalidate the overstated claim until corrected.	INVALID
CANDIDATE_GENERATOR_TREATED_AS_WARRANT	The system that generated a candidate action, plan, design, or output is treated as the admissibility warrant for that same candidate without independent evidence or anti-circularity disclosure.	Separate candidate-generation source from admissibility-warrant source, add independent evidence or anti-circularity justification, or refuse the physical/causal claim.	INVALID
METHODOLOGICAL_INTEGRITY_RECORD_MISSING	MFDI is triggered for a strong methodological, dependency, falsification, or necessity claim, but the required MethodologicalIntegrityRecord is absent.	Add the record, narrow the claim, label it exploratory, rerun affected surfaces, escalate, or refuse the stronger claim.	INVALID for affected strong claim
METHODOLOGICAL_STATUS_UNSUPPORTED	Declared methodological integrity status exceeds the evidence, dependency position, alternatives analysis, uncertainty basis, or revision-trigger record.	Downgrade status, add evidence, rerun dependent surfaces, or invalidate the affected methodological claim.	ESCALATE/ INVALID

Flag	Trigger	Required action	Severity
CROSS_DOMAIN_WARRANT_SUBSTITUTION	A scientific, engineering, legal-regulatory, normative, institutional, certification, compliance, commercial, or anthropomorphic label or warrant is used to support a materially different claim domain without a declared and supported bridge.	Record the claim type, claim domain, all material warrant domains, and one explicit bridge record for each material cross-domain relationship; narrow, escalate, redesign, or invalidate an unsupported cross-domain conclusion.	INVALID for affected strong claim
RIGHTS_RELEVANT_EFFECT_UNROUTED	A D6 Meaning or D7 Environment effect is rights-relevant or instantiates a protected interest under the Section 7 routing rule, but no corresponding rights-floor instance exists in an applicable covered dimension, no governed Section 12.8 extension is declared, and no reviewed non-materiality determination resolves the issue.	Construct and evidence the applicable rights-floor instance, declare a governed rights-coverage extension, or record a reviewed non-materiality determination; reopen RF/NCRC where material.	INVALID for affected rights claim

CBI PCC fields (Normative for triggered cases). When COMPUTABLE_BUT_INADMISSIBLE is triggered, the PCC MUST record: PCC.CBI.ComputedObject; PCC.CBI.InadmissibilityReason; PCC.CBI.Layer; PCC.CBI.AllowedUse; PCC.CBI.ProhibitedUse; PCC.CBI.ReviewStatus; and, where applicable, PCC.CBI.ReviewerNote or EscalationPath.

Stewardship audit-flag tokens are defined canonically in Appendix N.7 and MUST NOT be redefined in downstream specifications; downstream documents SHALL reference Appendix N.7.

No independent audit-flag mirror table is maintained in Appendix H.3. Implementers MUST consume Section 14.3A together with CANONICAL_AUDIT_FLAG_REGISTRY_v1.1.yaml, or a mechanically generated projection of that pair. Absence from an appendix, workbook, or companion does not retire a registry-listed controlling token.

Additional v9.0–v9.6 hardening flags (Normative). The canonical audit-flag set SHALL include PLSS_ESCALATION_TRIGGER_UNRESOLVED, DECISIVENESS_METHOD_POSTHOC, EK_NONDEFAULT_TIER2_INVALID, EK_PARAMETER_UNJUSTIFIED, UNCERTAINTY_PROVENANCE_MISSING, RIGHTS_CELL_UNKNOWN_GATE_CRITICAL, and CATASTROPHE_CELL_UNKNOWN_GATE_CRITICAL. These flags are triggered as specified in Sections 10.2A.7, 10.2A.8, 5.2, 5.7, and 15.1 and MUST be recorded in the PCC when applicable.

EK_NONDEFAULT_TIER2_INVALID trigger clarification (Normative). This legacy-stable token means that Tier 2 used non-default e_k without a valid, explicitly invoked named governed instrument in PCC.EKRegistry. It does not invalidate the named-instrument exception permitted by Section 5.2. The token name is retained for machine compatibility; implementations MUST use this trigger definition rather than infer that every Tier 2 non-default value is invalid.

EK_PARAMETER_UNJUSTIFIED (Normative). Trigger: non-default e_k is used without (i) a matching PCC.EKRegistry entry when required, (ii) declared rationale and source/instrument note, (iii) stated bounds, or (iv) the required tier-appropriate sensitivity disclosure. Required action: populate PCC.EKRegistry and missing disclosures, or reset e_k to 1.00 and recompute. Severity: INVALID.

14.3A.2 Representation, Evidence and Selection Token Completion (Normative)

This table consolidates tokens already required by the cited rules and vectors. Canon is their owner. Severity qualifies the affected claim; it does not replace the required action. NORMATIVE_UNDERDETERMINED is a NonDecisiveReason, not an additional gate or severity. Existing registry rows, including RLS_DEPENDENCE_SENSITIVE, remain controlling and are not duplicated here.

Token / type	Trigger	Required action and source
BASELINE_RANK_SENSITIVE Audit flag / ESCALATE	An admissible baseline comparison changes ordering or decisiveness.	Disclose both constructions; retain the gate-controlling floor reference; narrow/refuse unique residual selection pending governed resolution. [§10.2; R.28]
CONFIDENCE_RANK_SENSITIVE Audit flag / ESCALATE	The required adverse-confidence counterfactual changes a preferred option or decisiveness.	Disclose point and counterfactual results; prohibit unique framework selection while sensitivity remains unresolved. [§5.2C]
GAP_NUMERICAL_GUARD_SENSITIVE Audit flag / ESCALATE	Epsilon treatment changes decisive/non-decisive classification.	Disclose epsilon and results; refuse unique selection until treatment is governed independently of the preferred outcome. [§10.3A]
MASK_ASYMMETRY_INVALID Audit flag / INVALID	Unjustified option-specific masking or denominator manipulation changes the comparative construction.	Invalidate the comparative RLS claim; repair the common basis and rerun. Non-maskable protections remain active. [§10.2B; R.29]
NORMATIVE_UNDERDETERMINED NonDecisiveReason / not a severity	Legitimate ex-ante weight constitutions select different options after the evidence surface is sufficiently closed.	Retain REFUSE_DETERMINISTIC_SELECTION and this reason; record any separate accountable-authority choice. Do not inflate sigma to stand for normative disagreement. [§13; B.14A; R.31]

Token / type	Trigger	Required action and source
OPTION_SET_THIN Audit flag / REVIEW	Plausible required alternatives or meaningful challengers are omitted without adequate rationale.	Permit only bounded exploration; withhold strong unique-selection claims until option closure is reviewed. [§15.4B; R.32]
PARTITION_RANK_SENSITIVE Audit flag / ESCALATE	Admissible representation alternatives change the leader or decisiveness.	Retain all results and REFUSE_DETERMINISTIC_SELECTION unless representation is resolved independently of the desired result. [§5.5A; R.22]
PROBABILITY_ORDER_SENSITIVE Audit flag / ESCALATE	Evidence-supported selection-material likelihood/order alternatives change the leader or decisiveness.	Disclose and govern alternative results; no unconditional unique-selection claim. [§5.2B]
REFERENCE_RANK_SENSITIVE Audit flag / ESCALATE	Evidence-supported reference-class alternatives reverse the comparison.	Retain the reference-family alternatives and refuse reference-invariant unique selection until governed resolution. [§5.1A; R.27]
RIGHTS_PARTITION_SENSITIVE Audit flag / ESCALATE	Scope allocation or partitioning can hide the full unallocated rights harm.	Evaluate the unallocated token and resolve required rights fields before RF_PASS; ordinary welfare allocation cannot repair the rights claim. [§7.4.5]
RLS_ROBUSTNESS_INCOMPLETE Audit flag / ESCALATE	A triggered required robustness module is not evaluated or remains unresolved.	Complete the required record/evaluation; withhold ALLOW_FRAMEWORK_SELECTION while incomplete. [§10.4B; R.33]
TEMPORAL_PARTITION_SENSITIVE Audit flag / ESCALATE	Representation-only temporal re-tokenization changes the decision.	Restore the governed contiguous-window representation or resolve it independently; no unique selection while sensitive. [§5.3B; R.23]

Token / type	Trigger	Required action and source
MANDATORY_TAIL_CATEGORY_MISSING Audit flag / ESCALATE	A required tail family is absent without a governed omission disposition.	Repair the scenario library or establish permitted omission justification and rerun; no complete-TRC or unique-selection claim from the incomplete record. [§8.3.1; §14.3B; R.19]
NORMATIVE_INPUT_PROVENANCE_MISSING Audit flag / INVALID	A required normative-input register or material entry lacks provenance.	Reject the affected run as non-conformant until provenance is restored and affected calculations/claims are rerun. [§14.3B; Appendix L; R.11]
MPS_HYPOTHESIS_SENSITIVE Audit flag / ESCALATE	A governed welfare-inclusion hypothesis changes a material conclusion, preferred option or decisiveness.	Disclose affected stakeholders/hypotheses/results; collect evidence, choose an acceptable precautionary policy, narrow, preserve alternatives or refuse definitive ranking. No convenient intermediate coefficient. [§12.6; Appendix G; R.15]
RIGHTS_EFFECT_TOKEN_MISSING Audit flag / INVALID for the rights claim	Required rights-token identity is missing where allocation could affect the result.	Restore and evaluate the full unallocated token; RF_PASS is unavailable until all material fields are resolved. [§7.4.5]
MANDATORY_TAIL_CATEGORY_MISSING_WITH_JUSTIFICATION Audit flag / DISCLOSE	A mandatory-category omission has the justification required by the governing rule.	Record and review that justification and the scenario-discovery boundary; this flag alone establishes neither completeness nor a pass. [§8.3.1; §14.3B]
MANDATORY_TAIL_PROB_FLOOR_VIOLATION Audit flag / INVALID for TRC conformance	A required category's supported scenario mass does not meet the governed probability floor.	Repair probabilities/category coverage and rerun TRC; do not count overlapping categories as disjoint probability mass. [§8.3.1; §14.3B]

14.3B Minimal Validator Requirements (v10.6+; Normative)

Any conformant RippleLogic v13.0 validator MUST check the following:

- Every declared instance maps to ≥ 1 scope (InstanceMap completeness).
- Confirm PCC.ScopeCoverageDeclaration.Mode is declared as FULL_SCOPE or REDUCED_SCOPE; if Mode = REDUCED_SCOPE then ScopeCoverage block is present with a blind-spot statement and escalation triggers.
- If Mode = REDUCED_SCOPE: PCC.AdmissibilityAttestation.NCRC_full_Cr_computed is present and TRUE.

If Mode = REDUCED_SCOPE: PCC.AdmissibilityAttestation.TRC_full_Ccat_computed is present and (TRUE or NA).

If TRC is computed/required (i.e., not NA): verify mandatory tail categories are present (or justified) and per-category probability sums meet p_{floor} ; otherwise require the corresponding audit flags (MANDATORY_TAIL_CATEGORY_MISSING / ...WITH_JUSTIFICATION / MANDATORY_TAIL_PROB_FLOOR_VIOLATION).

If Mode = REDUCED_SCOPE: PCC.AdmissibilityAttestation.Containment_computed_as_applicable is present and (TRUE or NA).

If CONTAINMENT_UCI_UNAVAILABLE triggers: audit flag is present, and PCC.Containment.UCIUnavailableDisposition + PCC.Containment.UCIUnavailableRemediationTimeline are present (non-empty).

Tier-3 claim enforcement (Normative). If Implementation Tier = 3 and CONTAINMENT_UCI_UNAVAILABLE is present, the validator MUST reject Tier 3 compliance unless PCC.TierClaimDisposition = TIER_ATTEMPTED_NOT_ACHIEVED_DOWNGRADED and PCC.TierAchieved = 2.

- Non-maskable checks: rights-covered cells and catastrophe cells are unmasked (no bypass).
- If ALLOCATION used: $\sum \alpha \leq 1$ for each declared effect-token.
- SCI level meets tier minimum (Tier 2: $\text{SCI} \geq 1$; Tier 3: $\text{SCI} \geq 2$).
- If SCI equals tier minimum: Next-Run Upgrade action present.
- Tier 2+ uncertainty: $\sigma(u,d,a)$ declared for all active cells using Method A, B, or C.

If PCC.Uncertainty.Method = B: verify PCC.Uncertainty.CellConfidenceAggregationMethod is present and valid, and that $c(u,d,a)$ is recorded (or recomputable) for all active cells.

- Tier-2 mapping triggers: if any trigger holds, instance mapping is present.

Provenance binding (Normative). Any analyst-entered parameter that can steer admissibility or decisiveness - including uncertainty method choice, scenario probabilities p_s , PLSS prominence inputs q_u , and any non-default aggregation operator - MUST have a reason code, source note or elicitation note, and a compact sensitivity note if altered values plausibly change the outcome.

If any welfare-relevant state, rights-relevant state, or enabling condition materially affects admissibility, precautionary handling, selection, or tie-break outcome, PCC.NormativeInputRegister is present.

For each ADMISSIBILITY_MATERIAL, PRECAUTION_MATERIAL, SELECTION_MATERIAL, or TIEBREAK_MATERIAL entry in PCC.NormativeInputRegister, the validator MUST verify the presence of NormInputID, InputType, Materiality, AffectedObjects, EvidenceBasisType, and at least one of: EvidenceReference, GovernedInterpretationBasis (not NONE), or PrecautionaryBasis (not NONE).

If PrecautionaryBasis $\neq$ NONE, the validator MUST verify that the precautionary basis is explicitly stated, the affected object(s) are identified, and, for Tier 3, a ReviewerNote is present.

If OptionSymmetryStatus = OPTION_SPECIFIC_WITH_JUSTIFICATION, the validator MUST verify that a SensitivityNote is present.

If the register is required but missing, or if a required material entry is incomplete, the validator MUST set NORMATIVE_INPUT_PROVENANCE_MISSING and reject the run as non-conformant.

- Rights-covered and catastrophe cells: never bypassed by reduced-scope mode or masking.

If PLSS is declared: q_u is present for all scopes, the operator is declared, the PLSS weight vector sums to 1, and each w_u^{PLSS} is greater than or equal to its constitutional floor.

If UCI is used: every component is VALUE, ESTIMATE, UNKNOWN, or NA; NA components are excluded and weights renormalized; UNKNOWN components do not receive an inserted score. If U6 is material: every U6 record includes a valid U6_View_Tag and mixed values include decomposition.

If all $q_u = 0$ under PLSS: the validator MUST verify that the declared uniform residual fallback (or another declared fallback with justification) was applied.

If an alternative PLSS operator is used: justification and a sensitivity note are present in the PCC.

If PLSS is declared: an escalation-trigger review field is present and the PCC.PLSS block is present.

If a gate-critical UNKNOWN_IMPACT occurs in C_r or C_cat: the validator MUST reject any attempt to resolve admissibility by phantom-instance fallback alone and MUST require the corresponding audit flag plus a declared conservative bound or rerun / downgrade / escalation disposition consistent with Section 5.7.

If a non-maskable rights or catastrophe cell remains in RLS aggregation: the validator MUST verify floor-baseline consistency or the presence of a governed dual-baseline declaration that preserves the floor-reference value as the controlling gate-consistency baseline.

If a dual-baseline method is declared: the validator MUST verify that the PCC contains a per-cell or per-cell-class dual-baseline declaration identifying the affected welfare cells, the welfare baseline used, the gate baseline fixed at floor-reference, and the rationale for the exception.

14.3C Record Semantics, Display Integrity and Privacy (Normative)

A claim-bearing interface SHALL distinguish evidence/telemetry, EffectTokens, derived impact snapshots, gate/verdict records, and authorized intervention/outcome records. Bind them to run, option, baseline, schema, component map, configuration, units, sign convention, provenance, time horizon and evidence-status identifiers. Schema validity checks record structure, not truth, orthogonality, completeness of real-world discovery, rights passage or authority. The supplemental record schemas and scoped reference checks listed in the release appendix do not replace the full run validator or required domain evidence.

Use stable U1-U7/D1-D7 identifiers and named fields or an explicitly pinned vector order. Preserve unknown/contested/out-of-domain states separately from numeric zero. Timestamp fields distinguish observation time, receipt time, monotonic uptime and authorization validity. Numerical conversion requires finite inputs, range checks, a declared scale and rounding/overflow policy. Representation-only reordering, display rounding, duplicate transport and storage failover must not change the effect ledger or grant a new decision state.

Display the signed residual field and its contribution decomposition separately from raw system diagnostics, quality/freshness indicators, qualification states, authority and execution. Defaults may not imply PASS, live telemetry, empirical coherence, recovered welfare or a measured alignment percentage. Every displayed result must resolve to a recorded calculation or observation; synthetic demonstrations stay visibly synthetic.

A signature or hash binds specified data under its cryptographic assumptions; it does not prove evidence truth, welfare measurement, lawful authority or physical safety. A privacy-preserving proof establishes only its declared formal relation over committed inputs. Record the relation/circuit, public inputs, commitment and trust sources, setup/verification assumptions, freshness, replay protections and the unproven facts. Protect vulnerable people without requiring them to possess an account, secret key or resource reserve.

Preserve enough decision-material evidence for legitimate reconstruction, challenge, remedy and requalification while minimizing unnecessary sensitive retention. A hash of predictable personal data is not anonymization. Retention expiry, redaction, lawful holds and deletion require explicit policies; deletion of material evidence changes future claim support and must remain auditable without recreating the removed sensitive content.

14.3D Measurement-Maturity Validator Requirements (Normative)

For Tier 2-3 claim-bearing runs, a conformant validator MUST verify that the PCC includes all Appendix V measurement-maturity fields applicable to the run. At minimum, the validator MUST check:

PCC.UCI_Measurement_Readiness where UCI is used or unavailable at Containment or tie-break;

PCC.Kernel_Use_Posture where numerical propagation is used or explicitly omitted;

PCC.HDW_Capture_Audit for Tier 3 public or institutional weighting processes;

PCC.PLSS_Escalation_Register whenever PLSS is used; PCC.Deployment_Context_Profile for Tier 2-3 claim-bearing deployments; and PCC.Biological_Evaluation_Maturity where biological SGP evidence is invoked.

The validator **MUST** reject or downgrade any claim that treats UCI-M1, K0/K1, BEM-1/BEM-2, unreviewed HDW, or unchallenged PLSS localism as empirically validated, capture-resistant, or operationally complete. Missing maturity fields trigger PCC_MATURITY_SURFACE_INCOMPLETE. A maturity field marked provisional may support a worked-run or architecture claim, but **MUST NOT** support empirical validation, ProofPack readiness, Tier 4 readiness, validated UCI, validated kernel propagation, or biological measurement-completeness claims.

14.4 Five-Sentence Public Rationale (5SPR) (Normative)

Tier 2-3 PCC **MUST** include a Five-Sentence Public Rationale:

Element	Question
CONTEXT	What decision was made and why now?
OPTIONS	What options were considered?
CONSTRAINTS	What was eliminated by NCRC, TRC, and/or Containment (and why)?
SELECTION	Why the chosen option won among selectable options?
MONITORING	What follow-up will be tracked and when will NCAR Reflect occur?

14.5 Tier-4 Design Target Boundary Statement (Normative)

Tier 4 compliance claims are **PROHIBITED** in the current v13.0 release. Tier-4 content is design target only (Appendix I). No determinism, hash-bound replay, or ProofPack claims may be asserted until ProofPack is publicly replayable and independently verified.

14.6 Socio-Technical Embedding: Legitimacy, Contestation, and Institutional Use (Informative)

RippleLogic is designed for socio-technical governance settings where values are plural, power is unevenly distributed, and institutional incentives can distort decision quality. The framework therefore treats legitimacy as a procedural property rather than an assumed outcome.

Hybrid Democratic Weighting governs welfare weights through structured stakeholder participation with anti-capture safeguards, while constitutional rights and tail-risk constraints limit what can be traded away even when stakeholders disagree.

The Provenance and Compliance Certificate supports contestability: it makes normative parameters, scenario assumptions, subgroup handling, and constraint-gate outcomes explicit so that affected parties can challenge decisions with evidence rather than rhetoric.

Institutional embedding is expected to vary by context, but at minimum high-stakes deployments **SHOULD** include:

- Independent review capacity (at least one reviewer with no reporting relationship to the decision owner)
- Audit sampling (for example, random audit lotteries selecting $\geq 5\%$ of Tier 3 PCCs for independent review)
- An escalation pathway when rights, tail-risk, or containment conditions are near-binding or sensitive under perturbations
- Designated channels through which affected stakeholders can access the PCC (or the 5SPR at minimum) and submit evidence-based challenges
- A governance body or designated authority responsible for adjudicating challenges, with documented procedures and timelines

RippleLogic does not remove politics; it makes the structure of political and ethical disagreement legible and auditable. The framework's value proposition is not that it eliminates conflict but that it converts implicit, opaque tradeoffs into explicit, traceable, and challengeable ones.

Contestation architecture. The 5SPR (Five-Sentence Public Rationale) serves as the minimum contestability interface: any affected party should be able to read the 5SPR and understand what was decided, what was excluded and why, and what follow-up is planned. For Tier 3 decisions, the full PCC provides the detailed evidence base for substantive challenge. Challenges SHOULD be processed through a documented review procedure with:

- A defined response timeline (recommended: ≤ 30 days for initial response)
- Documented outcomes (challenge upheld, rejected with reasoning, or additional evidence requested)
- Version increment if the challenge results in parameter changes
- Preservation of the original PCC alongside the revised PCC

Power asymmetry awareness. In contexts where decision owners have significantly more power than affected stakeholders, the framework's structural protections (NCRC, worst-off subgroup semantics, mandatory tail categories, independent challenger requirements) serve as institutional counterweights. However, these protections are only effective if the institutional context supports their enforcement. Deployments in contexts with weak rule of law, captured institutions, or suppressed civil society SHOULD be flagged for enhanced scrutiny and may require external oversight partnerships.

14.7 Stewardship Layer (Normative): Applicability and Binding Semantics

14.7.1 Applicability Predicate (Normative)

RippleLogic Stewardship requirements apply only when a run is performed in a stewardship-relevant role (power asymmetry, delegated influence, public/institutional impact, or execution capability). Define:

$Stw_req(run)$ in $\{0,1\}$, where $Stw_req(run) = 1$ if and only if any of S1 through S4 holds. S5 is a modifier: it constrains provenance and influence controls when present, but S5 alone does not trigger $Stw_req(run)$.

The $Stw_req(run)$ predicate is defined in Section 14.7.1 and MUST NOT be redefined in downstream specifications; downstream documents SHALL reference this section.

- S1 (Delegated Influence): An assisting system is asked to recommend, rank, or narrow option space for a decision owner, such that the decision owner may reasonably rely on it.
- S2 (Public/Institutional Effect): Outputs are intended for public posting or institutional/policy use beyond the private decision owner.
- S3 (Execution Capability): The assisting system can execute actions (write/post/act) or trigger downstream actions beyond drafting.
- S4 (High-Stakes Rights Exposure): The assisted decision materially affects rights, safety, liberty, livelihood, or access to essential services, or produces irreversible/high-impact outcomes for protected stakeholders.

S5 (Authenticated Operator Channel; modifier): The system has an authenticated operator channel capable of altering configuration, permissions, or deployment posture; when S5 co-occurs with any of S1 through S4, additional provenance and influence disclosure requirements apply per Appendix N.

Interpretation notes (Normative).

S1 scope limiter: Delegated influence applies when the decision owner may reasonably rely on the system's output to materially shape a consequential decision (including meaningful option-narrowing or ranking that guides action). Purely informational or educational responses where independent judgment and verification capacity are preserved do not by themselves trigger S1.

S5 scope limiter (no loophole): S5 alone, absent any of S1 through S4, does not trigger Stewardship instrumentation. When S5 co-occurs with any of S1 through S4, $Stw_req(run)=1$ is triggered by the co-occurring condition(s) (S1–S4), not by S5.

If $\text{Stw_req}(\text{run}) = 0$, Stewardship instrumentation is OPTIONAL and Stewardship audit flags do not apply. If $\text{Stw_req}(\text{run}) = 1$, Stewardship instrumentation is REQUIRED per Appendix N.

14.7.2 Stewardship vs Cascade (Normative)

Stewardship does not modify the RippleLogic cascade definitions for NCRC, TRC, CSV, RLS, or UCI/HOI. Stewardship governs assistance integrity (boundary stability, authorship preservation, influence transparency). When Stewardship is required ($\text{Stw_req} = 1$), failure conditions are enforced via audit flags and PCC validity rules defined in Appendix N.

14.7.3 Flag Semantics Pointer (Normative)

Stewardship flags follow the canonical audit-flag severity semantics:

- INVALID flags invalidate the PCC and require rerun/remediation.
- REVIEW flags require mitigation and monitoring but do not by themselves invalidate the PCC.

Stewardship audit flags are defined in Appendix N and apply only when $\text{Stw_req}(\text{run}) = 1$.

14.8 Computational Closure Declaration (Normative)

Purpose. Every RippleLogic run SHALL declare the boundary of computation so that claims do not silently extend beyond the scopes, horizons, evidence classes, or propagation modes actually evaluated.

A PCC.ComputationalClosure object SHALL include at minimum:

- scopes_evaluated: Union Scopes included in the run;
- scopes_guardrail_only: Union Scopes protected through floors, NCRC, TRC, or CSV but not materially weighted in residual welfare ranking;
- scopes_outside_computation: scopes not evaluated beyond humility-horizon recognition;
- time_horizon_evaluated: declared decision horizon;
- time_horizon_outside_computation: time horizons acknowledged but not modeled;
- kernel_mode: NONE, QUICK, or declared domain-kernel-id;
- propagation_edges_included: declared kernel edges or NONE;
- propagation_edges_excluded: known plausible edges excluded from computation;
- externalities_excluded: known plausible externalities excluded, with reason codes;
- claim_boundary: exact wording of what the run may and may not claim.

A run SHALL NOT claim coverage over scopes, horizons, externalities, or propagation pathways outside its ComputationalClosure boundary. Any public-facing summary SHALL preserve this claim boundary.

SECTION 15: IMPLEMENTATION GUIDANCE AND TIERS (NORMATIVE)

15.0 Purpose (Normative)

This section specifies minimum compliance requirements by tier and provides operational guidance for applying RippleLogic in real decisions. Tiers define minimum obligations; deployments may exceed them. The framework is designed to improve over time via NCAR learning and governed updates.

15.1 Tier Requirements Matrix (Authoritative; Normative)

Normative authority: This matrix is the single authoritative statement of tier compliance. If any other sentence in this document conflicts with this matrix, this matrix governs.

Capability / Requirement	Tier 1 (Heuristic)	Tier 2 (Core, Strict TRC)	Tier 3 (Auditable)
Option set	MUST list ≥ 2 options	REQUIRED	REQUIRED
Baseline declaration	REQUIRED	REQUIRED	REQUIRED
Impact scale and Baseline-Zero	REQUIRED (qualitative allowed)	REQUIRED (quantitative recommended)	REQUIRED (quantitative plus auditable)
NCRC rights check	REQUIRED (heuristic minimum)	REQUIRED	REQUIRED
Worst-off subgroup for rights	Recommended	REQUIRED for directly affected rights cells	REQUIRED for rights-covered cells
TRC tail-risk	REQUIRED qualitative trigger screen; quantify upon escalation	REQUIRED when catastrophe relevance plausible	REQUIRED
Scenario set size when TRC used	N/A	≥ 5 minimum	≥ 20 minimum
Mandatory tail categories plus p_floor	Recommended if TRC used	REQUIRED when TRC used	REQUIRED
CSV qualification (Containment and Structural Viability)	REQUIRED before any ordinary ranking or selection claim; a brief qualitative screen or CSV_NOT_MATERIAL rationale is sufficient	REQUIRED proportional/assumption-bounded CSV result before RLS; quantitative UCI is not required unless material	REQUIRED full and binding
RLS scoring	Optional; if used, only after a pass-equivalent CSV result	REQUIRED only for selectable survivors after CSV	REQUIRED only for selectable survivors after CSV
Uncertainty plus discrimination band	Optional	REQUIRED (active cells)	REQUIRED
Tie-breaks (UCI/HOI)	Optional	Recommended	REQUIRED when RLS non-decisive
Kernel propagation	NONE default	NONE default; QUICK if KQS policy satisfied	QUICK only if KQS policy satisfied plus sensitivity; else NONE
KQS policy	Optional	REQUIRED if kernel used	REQUIRED if kernel used
Sensitivity analysis	Optional	Recommended	REQUIRED
PCC artifact	Optional	REQUIRED (basic)	REQUIRED (full)
Audit flags	Optional	REQUIRED when triggered	REQUIRED when triggered
Stewardship (conditional)	Optional	Required iff Stw_req=1	Required (Stw_req assumed TRUE for Tier 3)
Tier-4 claim	PROHIBITED	PROHIBITED	PROHIBITED
Stakeholder discovery and mapping (SDP/InstanceMap)	Optional	REQUIRED when any SDP trigger holds; otherwise a lightweight stakeholder note or voluntary SDP is RECOMMENDED	REQUIRED (InstanceMap with multi-scope redundancy handling)

Capability / Requirement	Tier 1 (Heuristic)	Tier 2 (Core, Strict TRC)	Tier 3 (Auditable)
			disclosures)
Stakeholder Coverage Index (SCI)	Optional	REQUIRED (SCI ≥ 1)	REQUIRED (SCI ≥ 2)

Tier 3 stewardship assumption (Normative): Any Tier 3 run SHALL set `Stw_req(run) = 1` by default because Tier 3 implies external scrutiny and governance-grade assistance posture.

15.2 Tier 1: Heuristic Application (“2-Minute Ripple Check”) (Normative minimum)

Use Tier 1 for low-stakes, reversible decisions where full calculation is disproportionate.

Tier 1 minimum protocol:

1. **Decision:** What must be decided, by when?
2. **Options:** List at least 2 options, including a third-path redesign option where possible.
3. **Rights screen (RF/NCRC heuristic):** Could any option plausibly violate LIFE, BODY, NEED, LBTY, DIGN, PROC, INFO, or ECOL? If yes, escalate to Tier 2+.
4. **Tail screen (TRC heuristic):** Could any option plausibly create catastrophic, irreversible, systemic, or ruin-path downside? If yes, escalate to Tier 2+.
5. **CSV screen:** Could the apparent local gain hollow out, overload, externalize onto, or depend on failure of a larger system, or does the option lack the resources, controls, authority, reversibility, or execution conditions needed to stand? Record `CSV_PASS_HEURISTIC`, `CSV_NOT_MATERIAL`, `CSV_REDESIGN_REQUIRED`, or escalate. An ordinary RippleLogic ranking or selection claim MUST NOT be made without a pass-equivalent CSV result. `CSV_PASS_HEURISTIC` and `CSV_PASS_WITH_CONTROLS_HEURISTIC` are human-facing Tier 1 teaching labels, not additional canonical machine tokens. A machine exporter must use the pinned interface and preserve Tier 1 plus heuristic evidence maturity; merely removing a suffix does not establish formal conformance or execution authority. Verified machine interchange requires the applicable registry/schema and replay evidence.
6. **Scopes touched:** Which operational scopes U1 through U7 are materially affected?
7. **Unioning move:** What redesign could reduce harms and increase shared benefit across the affected scopes?
8. **Choose and record:** Compare only options that passed the heuristic qualification screens. If no option passes, redesign, escalate, delay, or refuse rather than forcing a ranking.

Tier 1 recording (recommended): a short note or mini-PCC containing the eight answers and any `CSV_NOT_MATERIAL` rationale.

Tier-1 Worked Example (informative; shows Tier escalation):

Decision: Accept a new consulting project this month, or decline to protect time for household responsibilities.

Options: (A) Accept as proposed (10 hrs/week for 4 weeks). (B) Decline. (C) Redesign: accept only if scope is reduced to 4 hrs/week and meetings are time-boxed.

Rights screen (heuristic): No plausible LIFE/BODY/NEED/LBTY/DIGN violations. PROC/INFO: contract terms are clear, no coercion. PASS at Tier 1. If any need for deceptive data use or exploitative terms were suspected, escalate to Tier 2.

Tail screen (heuristic): No plausible catastrophic, irreversible, or systemic downside. PASS at Tier 1.

CSV screen: Option A plausibly creates a local income gain by hollowing household stability and personal recovery capacity; record `CSV_REDESIGN_REQUIRED`. Option B is structurally viable but loses the project benefit; record `CSV_PASS_HEURISTIC`. Option C includes bounded hours and time-boxed meetings; record `CSV_PASS_WITH_CONTROLS_HEURISTIC` if the limits are accepted and monitorable.

Scopes touched: U1 Self (Health, Agency, Meaning). U2 Household (Social, Material, Health). U3 Community/Organization impacts minor and reversible.

Unioning move: Option C reduces household harm while capturing some benefit. Compare B and the controlled form of C only after the CSV screen. Choose C if renegotiation succeeds and the controls are real; otherwise choose B if household strain would be material.

Record (mini-PCC): store the eight answers, CSV status, controls, and chosen option so NCAR learning can refine future time-budget decisions.

15.3 Tier 2: Core Calculable (Strict TRC and Proportional CSV Posture) (Normative)

Use Tier 2 for routine but consequential decisions requiring transparent computation and a basic PCC.

Tier 2 minimum obligations:

- Construct impacts on the [-1,+1] scale using the Section 5 pipeline.
- Run RF/NCRC with worst-off subgroup checks for directly affected rights cells, including categorical-prohibition and severe-rights-hazard channels where triggered.
- Execute TRC when catastrophe relevance is plausible: use at least 5 scenarios; include mandatory tail categories and probability floors unless implausible with documented justification; compute CVaR and the corridor check; and record the TRC table in the PCC.
- Complete a proportional, assumption-bounded CSV review for every option before ordinary RLS ranking. Tier 2 CSV MAY use qualitative or proxy evidence where validated structural instruments are unavailable, but it MUST issue one explicit status from the canonical ladder: CSV_PASS, CSV_PASS_WITH_CONTROLS, CSV_NOT_MATERIAL, CSV_REDESIGN_REQUIRED, CSV_ESCALATE, CSV_FAIL, or CSV_EMERGENCY_PROVISIONAL where the emergency protocol applies. A CSV_NOT_MATERIAL result MUST state why no material containment or structural-viability pathway is present. Refusal is a permitted claim or action disposition after CSV_FAIL, unresolved material uncertainty, or failed authority conditions; CSV_REFUSE is not a canonical CSV status token.
- Form the selectable set only from options with adequate RG claim support, RF/NCRC pass, TRC pass, and a pass-equivalent CSV status (CSV_PASS, CSV_PASS_WITH_CONTROLS, or CSV_NOT_MATERIAL).
- Compute RLS only for the selectable set and record the normalized ranking. If no option is selectable, redesign, delay, escalate, or refuse; do not force a score-based choice.
- If the selectable survivors are uncertainty-overlapped or non-decisive, apply only legitimate residual tie-break or authority-selection procedures; do not reopen or rescue a failed qualification level.
- Produce a basic PCC plus 5SPR, including the proportional CSV evidence, status, controls, blind spots, and review trigger.

Tier 2 propagation guidance:

- Default propagation_mode = NONE.
- QUICK propagation is allowed only if KQS policy is satisfied and sensitivity is feasible.

Structural-assurance note (Normative). Proportional CSV is mandatory before ordinary Tier 2 ranking, but its evidence depth is proportional to stakes and measurement maturity. A Tier 2 CSV result based on assumptions, proxies, or qualitative evidence MUST declare those limits. If the principal hazard is plausible hollowing-out, structural degradation, physical/causal inadmissibility, uncontrolled dependency, or other containment-relevant harm that cannot be responsibly bounded at Tier 2, the run MUST escalate to Tier 3, redesign, narrow the claim, or refuse the stronger selection claim. Quantitative UCI is not universally required at Tier 2; explicit structural reasoning and a pass-equivalent CSV status are.

15.4 Tier 3: Standard Auditable (Normative)

Use Tier 3 for high-stakes, contested, or externally scrutinized decisions where auditability and anti-gaming posture are required.

Tier 3 minimum obligations:

- Full PCC (Appendix H), including: impacts, subgroup semantics, scenario library, containment results, RLS, uncertainties, sensitivity, audit flags, signatures.
- NCRC: Subgroup analysis required for rights-covered cells.
- TRC: Required with ≥ 20 scenarios and mandatory tails plus probability floors (Appendix D).
- CSV / Containment and Structural Viability: Required and binding.
- Sensitivity analysis bundle required:
 - Weights perturbation,
 - Rights threshold perturbation (± 0.05 on θ_r as a sensitivity test),
 - Kernel perturbation if QUICK used,
 - Scenario probability perturbation,
 - Direct saturation β perturbation and, where RF/NCRC is material, rights saturation β_{RF} perturbation,
 - Baseline sensitivity for non-maskable FLOOR_REFERENCE cells when a defensible STATUS_QUO welfare counterfactual exists.
- Kernel: QUICK permitted only if QQS policy permits:
 - $QQS < 0.40$: NONE only,
 - $0.40-0.50$: QUICK plus mandatory sensitivity,
 - ≥ 0.50 : QUICK permitted plus sensitivity.
- FULL propagation is prohibited for Tier 1-3 claims.

15.4A Conformance, empirical readiness, and reference implementation posture (Normative clarification)

Tier 1-3 run-level conformance claims concern whether a particular run satisfies the requirements of this specification. They are distinct from framework-level empirical operational-readiness claims about comparative real-world performance, calibration quality, or superiority over alternative governance methods. Tier 3 run-level conformance means the run satisfies Tier 3 procedural, evidentiary, and audit requirements as declared; it does not by itself imply that every structural indicator family used for UCI or Containment is externally validated, mature across domains, or empirically calibrated beyond the declared measurement pack.

Reference implementation posture (Informative with binding disclosure consequence). This Foundation Paper normatively defines conformance without requiring a published reference implementation. However, project custodians SHOULD publish sibling reference implementation artifacts consisting at minimum of a machine-readable PCC schema, validator rule set, and executable Appendix R fixtures. Until such artifacts are published in a release manifest, framework-level claims of replay-grade operational readiness MUST remain provisional even where run-level Tier 1-3 conformance is possible.

Evidence-generating adoption posture (Normative research guidance). Where lawful, ethical, and operationally feasible, a preregistered sample of Tier 3 runs SHOULD include a shadow comparator, alternative method, or independent parallel review so ordinary use can contribute comparative evidence rather than only conformance records. The sampling fraction, comparator, outcome measures, and publication rule MUST be declared before results are known. A comparator is not required when it would create unacceptable rights, safety, confidentiality, or operational burdens; such exclusions MUST be recorded. This guidance creates no deployment permission and does not turn routine records into validation evidence without a declared study design. For comparative claims, declare the actual comparator rules and match evidence access, relevant rights/risk constraints and review resources. A simplified unconstrained scalar comparator supports only the corresponding ablation claim.

15.4B Option-Set Closure and Challenger Generation (Normative)

A ranking cannot establish superiority over a feasible option that was never generated. Tier 3 and high-stakes Tier 2 runs SHALL create an OptionClosureRecord before unique selection. The OptionClosureRecord MUST identify the option-generation methods and candidate sources used, the search or challenge effort proportionate to stakes,

who generated and challenged the alternatives, and why each material option class was included or excluded. Model-generated alternatives remain candidates whose feasibility, rights, risk, and viability require independent qualification.

The record SHALL consider, where feasible: no action or expected continuation; delay; staged or reversible trial; redesign; constitutive mitigation; materially cleaner or safer alternative; and a plausible hybrid/challenger option. Each omitted class requires a reason linked to feasibility, legality, evidence, or decision boundary.

An option set dominated by one favored proposal and deliberately weak comparators, or lacking a plausible no-action/redesign/challenger path without rationale, sets `OPTION_SET_THIN`. A thin set may support exploratory analysis but MUST NOT support a strong unique-selection claim until reviewed.

Option generation remains proportional. Low-stakes Tier 1 decisions may use a concise record; the rule is not a requirement to invent implausible options.

15.5 NCAR Integration by Tier (Normative)

All tiers SHOULD operate within NCAR.

Phase	Action
Notice	Define scope, unions, options, baseline, and configuration
Choose	Execute cascade and emit PCC
Act	Implement with monitoring aligned to predicted impacts and tail scenarios
Reflect	Compare observed outcomes to predictions; update indicators, kernels, scenario libraries, and weights through governed procedures

Reflect cadence defaults:

- Tier 2: Within 6 months or after major outcome data
- Tier 3: Within 3-6 months or after major outcome data
- Emergency/Fallback: Review cadence per declared severity and risk

15.6 Improvement and Future Upgrades (Normative stance)

RippleLogic is corrigible. Any improvement MUST:

- Preserve NCRC/TRC non-compensability structure,
- Be explicit, versioned, and auditable (AIL),
- Be tested under the validation program (Section 17),
- And be introduced via governed update processes (NCAR Reflect plus versioning).

SECTION 16: RELATIONSHIP TO EXISTING FRAMEWORKS (INFORMATIVE)

16.0 Purpose

This section positions RippleLogic relative to existing ethical, governance, and decision frameworks. It clarifies structural differences and interfaces.

16.1 Compared to Utilitarianism / Cost-Benefit Analysis (CBA)

CBA often scalarizes plural values into a single metric (money/utility), enabling rights tradeoffs. RippleLogic blocks this via NCRC (rights as constraints) and TRC (tail risk as constraint), then optimizes welfare only within the admissible set.

16.2 Compared to Deontology (Rule-based ethics)

Pure deontology can lack a complete operational procedure for comparing permitted options under uncertainty. RippleLogic retains non-compensable constraints but adds computable consequence modeling and tail-risk bounding.

16.3 Compared to Rawlsian Justice and Capability Approaches

Rawls provides priority of liberty and fair basic structure; capabilities provide multi-dimensional flourishing. RippleLogic operationalizes multi-dimensional welfare into a computable 7×7 matrix and makes rights floors explicit across unions.

16.4 Compared to MCDA (Multi-Criteria Decision Analysis)

Unconstrained compensatory MCDA models aggregate criteria without the non-compensable constraints used here (Keeney & Raiffa, 1976; Belton & Stewart, 2002). MCDA also includes constrained, lexicographic and non-compensatory methods; qualification before ranking is not unique to RippleLogic. Its comparison claim concerns the particular integration of rights, tail-risk and containment rules, residual RLS, uncertainty refusal and auditable interfaces. Comparative studies must name and implement a credible constrained comparator with matched evidence and resources; a pure weighted sum alone is an ablation, not evidence of superiority over MCDA generally.

16.5 Compared to AI Alignment Toolsets (RLHF, Constitutional AI, risk frameworks)

Many alignment approaches train systems toward proxies (human feedback, constitutional principles) without hard admissibility gates. RippleLogic provides a decision-engine architecture: action-space filtering by NCRC/TRC/CSV, structured welfare scoring (RLS), auditable traces (PCC), and corrigibility via NCAR.

16.6 Interoperability with Governance Standards

RippleLogic is designed to be interoperable with governance regimes emphasizing accountability and risk management (for example, NIST AI RMF) through its PCC record, scenario governance, and explicit risk bounding. It adds formal lexicographic rights and tail-risk operators that many standards leave at a principles level.

16.7 Compared to Commons Governance

RippleLogic is compatible with institutional approaches to governing shared resources and externalities. Where commons governance emphasizes rules-in-use, monitoring, graduated sanctions, and polycentric coordination, RippleLogic contributes a computable, auditable decision cascade that makes rights floors, tail-risk bounds, and cross-union ripple effects explicit in each decision record (Ostrom, 1990).

16.8 RLS Validation and Calibration Status (Informative bridge)

Current status. RLS is a candidate structured residual welfare-scoring specification and audit method, subject to the candidate acceptance checks and disclosed unresolved findings. It does not yet claim empirical validation of the dimensional non-overlap of the 7x7 welfare field, calibrated decisiveness thresholds, or inter-rater reliability of cell-filling across independent analysts. A companion RLS Validation Protocol should test these claims through scored-case datasets, correlation/factor analysis across the 49 cells, inter-rater reliability studies, disagreement heatmaps, scalar-comparator analysis, and revision rules. Until those studies are completed, public claims should say that RLS has an explicit auditable structure with disclosed specification and evidence limits, not that it is empirically proven to produce superior decisions.

SECTION 17: VALIDATION, FALSIFICATION, AND RESEARCH PROGRAM (NORMATIVE FOR CLAIMS)

17.0 Purpose (Normative)

RippleLogic makes testable claims about decision quality, rights protection, tail-risk avoidance, and auditability. This section specifies explicit falsification criteria and a staged validation program. Until such validation is completed, RippleLogic remains a theory-to-practice system with bounded claims: Tier compliance is claimable, real-world performance superiority is not.

This validation program follows a falsification-first posture consistent with scientific corrigibility: components that fail their empirical tests must be revised or abandoned rather than defended by authority (Popper, 1959).

17.0A Layered Falsification Attribution Rule (Normative for validation claims)

A failure must be attributed to the layer it actually falsifies before a broader falsification claim is made. A failed measurement does not by itself falsify the framework; a failed projection does not by itself falsify a rights floor; and a failed gate operator must not be excused as mere input error when all admissible data, projection, and parameter alternatives have been ruled out. A failed preregistered operational prescription remains a failure of that version and claim even when broader causal or operator attribution remains unresolved. Subsequent repaired versions do not retrospectively validate the failed prescription.

Every recorded validation failure SHOULD include both FalsificationClass and FalsificationLayer. Recommended FalsificationLayer values are DATA, PROJECTION, PARAMETER, GATE_OPERATOR, PRIMITIVE, GOVERNANCE, and COMMUNICATION.

Layer	What failure means	Does not by itself falsify	Required response
DATA	Input missing, wrong, stale, biased, or insufficient.	Gate definitions or framework primitives.	Re-collect, disclose uncertainty, rerun.
PROJECTION	Scenario, mapping, kernel, stakeholder model, or prediction surface was wrong.	Rights logic or the whole framework.	Freeze a new projection; rerun as a new run.
PARAMETER	Threshold, coefficient, weight, or corridor was mis-set or post-hoc.	Cascade structure.	Reset ex ante under governed procedure; rerun.
GATE_OPERATOR	NCRC, TRC, CSV, SVC, or RLS rule produces a reality-contradicted verdict across the preregistered finite or formally bounded admissible data/projection/parameter family, under its declared stopping and attribution rules.	A single bad input.	Challenge the operator definition; governed revision.
PRIMITIVE	A declared primitive cannot reconstruct a required phenomenon or produces unavoidable contradiction.	Any single failed run.	Reopen Primitive Audit.
GOVERNANCE	Authority, permission, legitimacy, contestability, or accountability failed.	Structural or mathematical cascade.	Governance review and authority correction.
COMMUNICATION	Public claim, summary, or disclosure exceeded claim boundary.	Underlying ethics or mechanics.	Correct the claim and update release/audit record.

Validator guidance. A failure logged without FalsificationLayer SHOULD trigger FALSIFICATION_LAYER_UNATTRIBUTED. A run that invokes PRIMITIVE or GATE_OPERATOR failure without first addressing plausible DATA, PROJECTION, and PARAMETER explanations SHOULD trigger

FALSIFICATION_LAYER_SKIPPED and require escalation. Before confirmatory testing, freeze the admissible attribution family, inclusion rules, stopping rule and treatment of unresolved alternatives. A formally bounded family needs a stated proof or exhaustive verification criterion. Keep failed locked tests visible; new alternatives or repairs require a new version/run and do not reopen the original success criterion. Where broader attribution cannot be resolved, report it as unresolved while withdrawing or narrowing the unsupported original performance claim.

Additional v12.0 do-not-collapse rules. Structural viability is not ethical permissibility. Stability is not justice. Permission is not admissibility. Output generation is not an admissible reference. A plan that can stand structurally can still violate rights; a generated answer can still lack a valid reference; and an authorized action can still fail the MathGov cascade.

17.1 Core Empirical Claims (Testable Hypotheses)

17.1A MNA-Layer Empirical Hypotheses (Normative research frontier)

H-MNA1 (model adequacy under stakeholder omission). Multi-agent predictive performance worsens when stakeholder welfare, rights, or enabling-condition variables are omitted from world models in settings where cooperation, conflict, or harm depend materially on them.

H-MNA2 (instability under normative omission). Architectures that systematically ignore stakeholders' welfare conditions, rights floors, or enabling conditions exhibit higher instability, defection, tail risk, or containment failure in repeated multi-agent environments than otherwise comparable architectures that track them explicitly.

H-MNA3 (cascade advantage over unconstrained scalarization). Rights-first and catastrophe-bounded cascade architectures reduce ruin-like outcomes, rights breaches, or coherence losses relative to unconstrained scalar optimization in comparable decision classes.

Hypothesis	Description
H1	Rights coherence: Decisions that pass NCRC produce fewer rights infringements than comparable baseline decisions, controlling for context.
H2	Tail-risk effectiveness: Decisions constrained by TRC exhibit lower realized tail losses than comparable decisions without TRC.
H3	Ripple exact sign accuracy (provisional research hypothesis): Test whether, after NCAR updates, ExactSignAccuracy reaches 0.70 on a preregistered evaluation set. This retained hypothesis is not a validated acceptance threshold or an automatic transfer of the legacy partial-credit cutoff. Appendix J.3 requires metric-specific justification, no-change and majority baselines, class results and uncertainty before confirmatory skill or readiness claims.
H4	Structural early warning: Persistent HOI > 0 predicts subsequent structural degradation better than baseline KPI monitoring.
H5	Anti-gaming effectiveness: In adversarial tests, specification gaming succeeds less often ($\geq 30\%$ reduction) relative to comparable governance processes without PCC plus mandatory tails plus subgroup semantics.

Because ripple effects partially track cooperation and network propagation dynamics, validation should also test whether kernel updates improve predictive performance in contexts where cooperation and trust dynamics are known to matter (Axelrod, 1984; Nowak, 2006; Christakis & Fowler, 2009).

H6 (PLSS prominence convergence). For recurring decision classes with standardized guidance, independently trained evaluators applying the same evidence base will produce PLSS prominence vectors with moderate or better convergence, operationalized in pilot testing by a pre-declared inter-rater agreement threshold.

H7 (PLSS floor activation). In a non-trivial fraction of PLSS-assessed decisions, constitutional floor weights assigned to guardrail scopes will contribute materially to final welfare ranking differences, demonstrating that guardrail scope retention is operational rather than merely decorative.

H8 (PLSS escalation fidelity). In pilot validation, plausible ET-trigger conditions will be detected and logged with sufficient sensitivity that PLSS posture is not systematically retained in clearly non-local cases.

17.1B Taxonomy and Implementation Falsification

Claims of improved dimensional separation, irreducibility or practical superiority require preregistered comparisons against plausible merges, splits and simpler alternatives. Report effect-boundary confusion, duplicate-effect coding, omitted outcomes, subgroup visibility, classification agreement, burden, explanation quality and decision reversals separately. Any composite of these criteria requires a justified scale and ex-ante weighting; adding incomparable loss quantities does not prove an optimum. Causal covariance is not an automatic redundancy penalty.

The adversarial fixture suite shall include sign inversions, missing/unknown values, duplicate effects/messages, evidence revisions, kernel direction and double propagation, nonlinear attribution, late or ungrounded controls, resource depletion, repeated recovery, uncertain outcomes, isolated operation, authority/quorum failure, privacy-proof overclaims and source/render drift. Synthetic scenario values and engineering stress-test traces test implementation or reasoning under assumptions; they are not observed societal data or empirical validation of the framework.

Use component-specific evidence labels: defined, implemented within named scope, tested under named fixtures, empirically validated for a domain, and independently validated are distinct. A passing unit test, rendered dashboard, replicated database, cryptographic proof or compiled artifact cannot upgrade a different claim. The exact verification report governs what was tested in this package.

17.2 Falsification Criteria (Normative)

17.2A MNA-Layer Falsifiers (Normative)

F-MNA1. Omission of welfare, rights, or enabling-condition variables does not reduce predictive adequacy in relevant multi-agent domains where those variables were hypothesized to matter materially.

F-MNA2. Architectures ignoring such conditions do not exhibit higher instability, tail risk, rights breach incidence, or containment failure than comparable architectures that track them explicitly.

F-MNA3. Cascade-constrained policies do not outperform unconstrained scalarized policies on declared rights, coherence, or tail-risk metrics in the relevant validation class.

Defect attribution prior to invoking F1-F5 follows the Appendix AE.3 layer matrix: each failure should be assigned to the layer it actually falsifies before any broader empirical falsifier is invoked.

RippleLogic components must be revised (or rejected) if evidence meets any of the following:

Criterion	Description
F1	NCRC failure: NCRC-passing decisions systematically produce worse rights outcomes than NCRC-failing decisions.
F2	TRC failure: TRC-constrained decisions show no reduction in realized tail losses compared to controls.
F3	Ripple predictiveness failure: the legacy < 60% sign-accuracy threshold is a provisional research falsifier; metric-specific adoption under J.3 is required before application across successive NCAR cycles for key cells.
F4	Structural safeguard failure: UCI/HOI do not correlate with or predict meaningful degradation.
F5	Anti-gaming failure: Red-team exercises repeatedly exploit predictable loopholes (> 30% of adversarial runs).

17.3 Validation Phases (Normative roadmap)

Phase 1: Formal verification and implementation testing (0-6 months)

- Independent implementation of Tier 2-3 algorithms from spec.
- Unit tests for canonical equations (Appendix B).
- Adversarial tests for masking, subgroup erasure, scenario omission, confidence inflation.
- Produce reference PCCs and reproducibility checks.

Phase 2: Measurement validation (4-15 months)

- Indicator reliability and validity for welfare dimensions and UCI components.
- Cross-population measurement checks where relevant.
- Calibrate magnitude anchors and uncertainty mappings.

Phase 3: Controlled pilots (10-24 months)

- Pilot deployments in organizations or municipalities.
- Compare RippleLogic vs baseline governance processes on: rights incidents, tail losses, stakeholder legitimacy, audit completeness, decision reversal rates.

Phase 4: Field studies and scaling (18-36+ months)

- Longitudinal monitoring of UCI/HOI and realized outcomes.
- Kernel calibration and scenario library refinement.
- Comparative performance across domains.

17.4 Metrics and Targets (Normative defaults)

17.4A Inter-Rater Reliability Commitment (Normative commitment)

IRR single-source rule (Normative).

The target table in this section controls. The RLS Validation Protocol mirrors and operationalizes it; any divergent threshold or statistic in a companion is non-controlling until adopted here. The numeric targets below are retained as provisional research targets. They become a reliability acceptance criterion only after a design-specific mapping is preregistered and its applicability justified; they are not interchangeable validated cutoffs for different statistics.

Before Tier 3 operational-readiness claims are made, the framework SHALL publish or register at least one inter-rater reliability (IRR) study on a representative case set covering: (i) welfare-cell anchoring, including magnitude construction μ_k using both percentile anchoring and threshold anchoring methods, (ii) rights-floor evaluation, including subgroup identification and worst-off impact estimation, (iii) local-scope/PLSS usage, including prominence signal construction and scope classification, and (iv) scenario construction and probability assignment for TRC.

The study MUST report the case set, evaluator instructions, adjudication procedure, disagreement classes, and reliability statistics appropriate to the data type used. For each component, freeze the measurement level and unit, case/rater design, primary preregistered reliability statistic and exact variant (including ICC model, unit and agreement/consistency choice where applicable), applicable target, point-estimate or confidence-interval decision rule, prevalence/class distribution, dependence and missingness treatment, and undefined-statistic handling. The companion protocol supplies design guidance, not permission to choose a favorable statistic after observing results. Raw percent agreement cannot substitute for an ICC, kappa or other named criterion. Undefined results are unresolved; they cannot demonstrate readiness. Until this mapping is adopted, reliability-based operational-readiness claims remain unavailable.

Component	Provisional target	Provisional minimum	Notes
Magnitude construction (μ_k)	≥ 0.70	≥ 0.60	Tested separately for percentile and threshold anchoring
Rights-floor pass/fail (NCRC)	≥ 0.80	≥ 0.70	Binary gate; higher target reflects non-compensability

Component	Provisional target	Provisional minimum	Notes
TRC pass/fail	≥ 0.80	≥ 0.70	Binary gate
RLS ranking order	≥ 0.70	≥ 0.60	Ordinal agreement on option ranking
PLSS prominence classification	≥ 0.70	≥ 0.60	Prominent vs Guardrail per scope

Sustained inability to exceed the minimum acceptable threshold on any component after rubric training indicates that the component requires redesign or additional anchoring guidance. Apply that interpretation only to the preregistered estimator/design mapping; do not infer measurement failure by moving an unvalidated numerical cutoff between statistics.

Until such evidence exists, Tier 3 remains architecturally specified but empirically incomplete.

Metric	Definition	Default Target
Rights violation rate	Verified post-decision rights infringements per protected subgroup	Lower than baseline governance
Tail-loss severity	Realized catastrophe-cell loss in worst outcomes	Lower than baseline governance
Exact sign accuracy	ExactSignAccuracy: exact predicted/observed sign agreement	0.70 hypothesis; metric-specific adoption required (J.3)
Magnitude error (RMSE)	RMSE on normalized $[-1,+1]$ for measurable cells	≤ 0.25 (domain-dependent)
Audit completeness	Proportion of required PCC fields correctly filled	≥ 0.95 Tier 3
Gaming success rate	Fraction of adversarial attempts that change outcome illegitimately	Decreasing over time; investigate if > 0.30

17.5 Study Design Requirements (Normative)

- Pre-registration required for pilots claiming performance improvements.
- For causal or comparative-effect claims, a claim-specific identification plan is REQUIRED: target estimand; intervention and credible comparator/counterfactual; population and horizon; identification assumptions; selection, confounding and interference treatment; measurement; and an explicit identification-status verdict. The design may be experimental or a justified non-experimental design. Pre/post change alone does not identify an incremental effect. Without adequate causal warrant, restrict claims to the supported descriptive, predictive or procedural scope. A comparator remains recommended for other exploratory performance work.
- Blinding recommended for backtests: analyst should not know outcome during the run.

17.5A Projection Pre-Registration

17.5A.5 Freeze-Hash and Unique-Prediction Requirement (Normative for public validation and ProofPack-prep claims)

A validation claim is not falsifiable unless the projection was frozen before outcome evidence was consulted. For Tier 3 public pilots, public validation claims, and ProofPack-prep runs, ProjectionPreRegistration SHOULD include projection_freeze_hash, freeze_timestamp_utc, prediction_scope, predicted_observable_consequence_class, outcome_observation_timestamp, and post_observation_modification.

A material validation claim SHOULD produce exactly one predicted observable consequence class or a declared finite prediction set with pre-specified resolution rule. If the projection produces no prediction, an

open-ended prediction, or a non-unique prediction without resolution rule, the claim is exploratory rather than falsifiable and SHOULD be marked NON_UNIQUE_PREDICTION or EXPLORATORY_ONLY.

If post_observation_modification=TRUE for a material validation claim, the run MUST NOT be used as a falsification test for that claim and SHOULD set POST_HOC_PROJECTION_INVALID with severity INVALID. A revised projection requires a new freeze hash and a new run. (Normative for pilots and Tier 3; Recommended otherwise)

Purpose. To preserve falsifiability and prevent post-hoc projection drift, pilot runs and Tier 3 runs SHALL pre-register the projection map from declared inputs to admissibility verdict before decision data is consulted in a way that could influence parameter selection.

17.5A.1 ProjectionPreRegistration object (Normative)

A ProjectionPreRegistration object SHALL be emitted, signed, timestamped, and stored before any cascade computation that contributes to a pilot or Tier 3 decision record. It SHALL contain at minimum: decision boundary, time horizon, and option set; stakeholder discovery record, including LSS and instance-to-scope mapping; rights coverage set C_r and per-right thresholds theta_r, including any non-default values with justification; catastrophe cell set C_cat, scenario set S, scenario probabilities p_s, CVaR confidence alpha, and TRC corridor tau_TRC; Containment trigger conditions and containing scopes; constitutional weight floors and any tier-appropriate non-default weights with justification; SGP v8.8 binding for impacted stakeholders, including target/record id, MPS band and interval, FPP status, validity status, and planned interval-sensitivity points; subgroup uplift gamma_subgroup and any non-default value with justification; kernel mode and kernel quality status where applicable; sensitivity sweep specification; pre-declared falsification, refusal, escalation, and re-run conditions.

17.5A.2 Modification discipline (Normative)

Once emitted, the ProjectionPreRegistration object SHALL NOT be materially modified after observation in any way that affects admissibility, selectability, selection, or conformance claims.

If material modification is required, the current run SHALL enter Refusal under Section 4.9, document the reason in PCC.RefusalRecord, and emit a new ProjectionPreRegistration object before any subsequent run.

Non-material corrections, including typographical corrections, reference updates, or formatting changes that do not alter values or verdict logic, MAY be applied with version increment and documented diff.

17.5A.3 Tier application (Normative)

Tier 3: Projection Pre-Registration is REQUIRED for all runs.

Tier 2: Projection Pre-Registration is REQUIRED for pilots and RECOMMENDED otherwise.

Tier 1: Projection Pre-Registration is RECOMMENDED where stakes are medium or high.

17.5A.4 Storage and audit (Normative)

The ProjectionPreRegistration object SHALL be stored with the PCC and referenced by hash. An independent reviewer SHALL be able to reconstruct the declared projection map from the stored object using only the applicable canon, declared evidence base, and referenced companion artifacts.

17.6 Backtesting Protocol (Normative)

A backtest run MUST:

- Reconstruct the information available at decision time (no hindsight leakage),
- Define an option set consistent with the historical context,
- Execute RippleLogic using only those inputs,
- Record a "Backtest PCC" with full trace,

- Score predictions against realized outcomes using declared metrics.

17.7 Open Science and Responsible Disclosure (Normative intent)

Where feasible, publish:

- Anonymized PCC datasets,
- Scenario libraries and kernels (with evidence notes),
- Failure cases and revisions.

Security note: Where publishing would create exploitation risk, disclosure should be responsible and staged, but internal auditability must be preserved.

SECTION 18: APPLICATIONS AND USE CASES (INFORMATIVE)

18.1 AI Deployment Governance (Tier 3)

Decision context: An organization is deciding whether to deploy an AI system that affects users at scale.

Key obligations:

- NCRC: Protect dignity, information integrity, due process for worst-off subgroups.
- TRC: Model catastrophic tails (infrastructure disruption, mass manipulation, cascading conflict).
- Containment: Avoid organizational gains that degrade Polity/Humanity coherence (legitimacy collapse, epistemic fragmentation).
- PCC: Produce auditable trace; ensure comparability across options.

18.2 Climate and Energy Policy (Tier 3)

Decision context: A polity evaluates energy transition policies that impact households, industries, and biosphere stability.

Key obligations:

- NCRC: Protect basic needs (energy poverty), health, life.
- TRC: Include climate tipping cascade scenarios; enforce corridor.
- Containment: Prevent short-term gains from degrading Biosphere UCI beyond tolerance.

18.3 Organizational Strategy (Tier 2-3)

Decision context: Organization deciding on major restructuring, automation, supply chain change.

Key obligations:

- NCRC: Protect basic needs and dignity for worst-off employees/communities.
- TRC: Required at Tier 2 if catastrophe relevance is plausible.
- Containment: Detect hollowing-out risk even if near-term welfare improves.

18.4 Personal and Household Decisions (Tier 1)

Decision context: Personal choices (job change, relocation, major purchase) where stakes are limited and reversible.

Use Tier 1 heuristic: options, rights screen, tail screen, unions touched, unioning redesign.

Escalate to Tier 2 if any plausible rights risk exists or if the decision touches catastrophe relevance.

SECTION 19: LIMITATIONS AND NON-TARGETS (NORMATIVE BOUNDARIES)

19.1 Epistemic Limitations (Normative acknowledgment)

Input quality dependence: RippleLogic outputs are only as good as the evidence and estimation process. The framework reduces predictable failures but does not guarantee correct forecasts.

Measurement difficulty: Some dimensions (Meaning; aspects of Agency) are harder to measure reliably across cultures. RippleLogic treats such measurement with caution and uncertainty documentation.

Kernel uncertainty and operational status: Ripple propagation via the QUICK kernel mode is architecturally specified and canon-complete for Tier 1-3 implementation. However, the Starter KOPS (Appendix K) carries evidence classes B and E with a global shrink factor of 0.35, and the KQS policy (Section 6.5) requires KQS ≥ 0.40 for any kernel use. In practice, most early Tier 2-3 runs are expected to operate in NONE propagation mode because domain-specific kernels with sufficient evidence coverage, identifiability, and predictive validation do not yet exist. The NONE mode is a valid and conformant propagation posture, not a workaround. Building empirically grounded kernels with KQS ≥ 0.50 is a Phase 2-3 validation priority (Section 17.3). Ripple-awareness is already operational in the current release line through the multi-scope welfare matrix (which makes cross-scale consequences visible), through containment (which detects cross-scale degradation), and through UCI/HOI (which monitors structural integrity). Numerical kernel-based propagation remains a design-ready capability awaiting domain-specific calibration. The current release line mitigates kernel uncertainty with KQS policy, sensitivity requirements, and humility fallback to NONE.

Scenario incompleteness: Scenario libraries cannot enumerate all tail risks; mandatory tail categories reduce omission but cannot eliminate deep uncertainty.

19.2 Institutional and Governance Prerequisites

Tier 3 requires trained analysts and governance capacity: subgroup analysis, scenario governance, structural indicators for UCI, and audit review.

HDW governance can be captured if safeguards are weak. RippleLogic mandates anti-capture mechanisms, but real-world institutions must implement them.

19.3 Risks of Compliance Theater

Any framework can be used as a rubber stamp. RippleLogic mitigates via:

- Hard gates (NCRC/TRC/CSV),
- Non-maskable cells and audit flags,
- Challenger requirement in emergency mode,
- Required PCC traceability and sensitivity.

Nonetheless, compliance theater remains a risk; validation and external scrutiny matter.

19.4 Non-Targets (Normative)

RippleLogic is NOT designed to:

- Replace personal moral dialogue in intimate relationships,
- Dictate aesthetic or spiritual preferences,
- Resolve all deep metaphysical disagreements,
- Provide certainty under irreducible uncertainty,
- Function as a rhetorical weapon to end debate.

It is a public decision operating system: a method for auditable, multi-stakeholder, high-consequence decisions.

SECTION 20: CONCLUSION: PROSPEROUS FUTURES FOR INTELLIGENCES

20.1 Prosperity Defined (Normative operational definition)

In RippleLogic terms:

Prosperity = rights-safe (NCRC) + tail-safe (TRC) + containment-safe + net-positive welfare across unions + coherence preserved (UCI/HOI safeguards) + corrigible learning (NCAR)

This definition blocks false prosperity:

- “Growth” that violates rights is not prosperity.
- “Efficiency” that increases catastrophic exposure is not prosperity.
- “Success” that hollows out structural coherence is not prosperity.

20.2 RippleLogic Thesis

Alignment across humans, institutions, and AI requires non-compensable constraints on rights and catastrophic risk, explicit modeling of ripple propagation through layered and intersecting analytical scopes, transparent and auditable decision procedures, and corrigible learning loops. These requirements are motivated by interdependence in complex systems (Meadows, 2008; Newman, 2010), planetary life-support constraints (Steffen et al., 2015; Rockström et al., 2009, 2023), and the asymmetric tail-risk properties of ruin dynamics (Taleb, 2012; Ord, 2020).

20.3 Call to Test, Critique, and Deploy (Normative stance)

RippleLogic v13.0 is a spec-hardened scaffold intended to be tested, falsified where wrong, refined where incomplete, and used proportionally to stakes. The method is designed to move societies toward ethically, sustainably, harmoniously, prosperously, and verifiably coordinated flourishing by making decision structure explicit, auditable, and corrigible.

REFERENCES (APA 7)

- Amodei, D., Olah, C., Steinhardt, J., Christiano, P., Schulman, J., & Mané, D. (2016). Concrete problems in AI safety. arXiv. <https://doi.org/10.48550/arXiv.1606.06565>
- Arrow, K. J. (1963). Social choice and individual values (2nd ed.). Yale University Press.
- Artzner, P., Delbaen, F., Eber, J.-M., & Heath, D. (1999). Coherent measures of risk. Mathematical Finance, 9(3), 203–228. <https://doi.org/10.1111/1467-9965.00068>
- Axelrod, R. (1984). The evolution of cooperation. Basic Books.
- Barabási, A.-L., & Albert, R. (1999). Emergence of scaling in random networks. Science, 286(5439), 509–512. <https://doi.org/10.1126/science.286.5439.509>
- Becker, G. S. (1981). A treatise on the family. Harvard University Press.
- Belton, V., & Stewart, T. J. (2002). Multiple criteria decision analysis: An integrated approach. Kluwer Academic Publishers.
- Bostrom, N. (2014). Superintelligence: Paths, dangers, strategies. Oxford University Press.
- Bowlby, J. (1988). A secure base: Parent-child attachment and healthy human development. Basic Books.
- Christakis, N. A., & Fowler, J. H. (2009). Connected: The surprising power of our social networks and how they shape our lives. Little, Brown and Company.
- Council of Europe. (1950). Convention for the Protection of Human Rights and Fundamental Freedoms (European Convention on Human Rights). <https://www.echr.coe.int/the-convention>

- Dunbar, R. I. M. (1992). Neocortex size as a constraint on group size in primates. *Journal of Human Evolution*, 22(6), 469–493. [https://doi.org/10.1016/0047-2484\(92\)90081-J](https://doi.org/10.1016/0047-2484(92)90081-J)
- Dunbar, R. I. M. (1993). Coevolution of neocortical size, group size and language in humans. *Behavioral and Brain Sciences*, 16(4), 681–735. <https://doi.org/10.1017/S0140525X00032325>
- Edmondson, A. (1999). Psychological safety and learning behavior in work teams. *Administrative Science Quarterly*, 44(2), 350–383. <https://doi.org/10.2307/2666999>
- Estes, J. A., Terborgh, J., Brashares, J. S., Power, M. E., Berger, J., Bond, W. J., Carpenter, S. R., Essington, T. E., Holt, R. D., Jackson, J. B. C., Marquis, R. J., Oksanen, L., Oksanen, T., Paine, R. T., Pickett, E. K., Ripple, W. J., Sandin, S. A., Scheffer, M., Schoener, T. W., ... Wardle, D. A. (2011). Trophic downgrading of planet Earth. *Science*, 333(6040), 301–306. <https://doi.org/10.1126/science.1205106>
- Floridi, L. (2013). *The ethics of information*. Oxford University Press.
- Goodhart, C. A. E. (1984). Problems of monetary management: The UK experience. In C. A. E. Goodhart (Ed.), *Monetary theory and practice: The UK experience* (pp. 91–121). Macmillan.
- Hammond, J. S., Keeney, R. L., & Raiffa, H. (1998). *Smart choices: A practical guide to making better decisions*. Harvard Business School Press.
- Integrated Food Security Phase Classification (IPC) Global Partners. (n.d.). Integrated Food Security Phase Classification (IPC) framework and tools. Retrieved February 19, 2026, from <https://www.ipcinfo.org/>
- Intergovernmental Panel on Climate Change. (2023). *Climate change 2023: Synthesis report*. <https://doi.org/10.59327/IPCC/AR6-9789291691647>
- International Labour Organization. (1998). ILO Declaration on Fundamental Principles and Rights at Work and its Follow-up. Retrieved February 19, 2026, from <https://www.ilo.org/ilo-declaration-fundamental-principles-and-rights-work/about-declaration/text-declaration-and-its-follow>
- International Labour Organization. (2012). ILO indicators of forced labour (Special Action Programme to Combat Forced Labour). <https://www.ilo.org/publications/ilo-indicators-forced-labour>
- Kahneman, D., & Tversky, A. (1979). Prospect theory: An analysis of decision under risk. *Econometrica*, 47(2), 263–291. <https://doi.org/10.2307/1914185>
- Keeney, R. L., & Raiffa, H. (1976). *Decisions with multiple objectives: Preferences and value tradeoffs*. Wiley.
- Lenton, T. M., Held, H., Kriegler, E., Hall, J. W., Lucht, W., Rahmstorf, S., & Schellnhuber, H. J. (2008). Tipping elements in the Earth's climate system. *Proceedings of the National Academy of Sciences*, 105(6), 1786–1793. <https://doi.org/10.1073/pnas.0705414105>
- Manheim, D., & Garrabrant, S. (2018). Categorizing variants of Goodhart's law. *arXiv*. <https://doi.org/10.48550/arXiv.1803.04585>
- March, J. G., & Simon, H. A. (1958). *Organizations*. Wiley.
- Max-Neef, M. A. (1991). *Human scale development: Conception, application and further reflections*. Apex Press.
- McGaughran, J. (2026). *Sentience Gradient Protocol (SGP) v8.8 [Specification]*. MathGov Institute for Ethical Systems Design. Historical lineage: SGP v8.5 [Specification]; SGP v8.4 [Specification]; Sentience Gradient Protocol (SGP) v7.0 [Specification]; v6.0 [Specification]; v5.6 [Specification]; v5.5 [Specification]; v5.3 [Specification]; v5.2.1 [Specification]; v5.2 [Specification]; and v4.7.2 [Historical predecessor specification].
- Meadows, D. H. (2008). *Thinking in systems: A primer*. Chelsea Green Publishing.
- National Institute of Standards and Technology. (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0) (NIST AI 100-1)*. <https://doi.org/10.6028/NIST.AI.100-1>
- Newman, M. E. J. (2010). *Networks: An introduction*. Oxford University Press.

- North, D. C. (1990). *Institutions, institutional change and economic performance*. Cambridge University Press.
- Nowak, M. A. (2006). Five rules for the evolution of cooperation. *Science*, 314(5805), 1560–1563. <https://doi.org/10.1126/science.1133755>
- Nussbaum, M. C. (2011). *Creating capabilities: The human development approach*. Harvard University Press.
- Odum, E. P. (1971). *Fundamentals of ecology* (3rd ed.). W. B. Saunders.
- Ord, T. (2020). *The precipice: Existential risk and the future of humanity*. Hachette Books.
- Ostrom, E. (1990). *Governing the commons: The evolution of institutions for collective action*. Cambridge University Press.
- Popper, K. (1959). *The logic of scientific discovery*. Hutchinson.
- Putnam, R. D. (2000). *Bowling alone: The collapse and revival of American community*. Simon & Schuster.
- Rawls, J. (1971). *A theory of justice*. Harvard University Press.
- Reporters Without Borders. (n.d.). World Press Freedom Index. Retrieved February 19, 2026, from <https://rsf.org/en/index>
- Rockafellar, R. T., & Uryasev, S. (2000). Optimization of conditional value-at-risk. *Journal of Risk*, 2(3), 21–41. <https://doi.org/10.21314/JOR.2000.038>
- Rockström, J., Gupta, J., Qin, D., Lade, S. J., Abrams, J. F., Andersen, L. S., Armstrong McKay, D. I., Bai, X., Bala, G., Bunn, S. E., Ciobanu, D., DeClerck, F., Ebi, K., Gifford, L., Gordon, C., Hasan, S., Kanie, N., Lenton, T. M., Loriani, S., ... Zhang, X. (2023). Safe and just Earth system boundaries. *Nature*, 619(7968), 102–111. <https://doi.org/10.1038/s41586-023-06083-8>
- Rockström, J., Steffen, W., Noone, K., Persson, Å., Chapin, F. S., Lambin, E. F., Lenton, T. M., Scheffer, M., Folke, C., Schellnhuber, H. J., Nykvist, B., de Wit, C. A., Hughes, T., van der Leeuw, S., Rodhe, H., Sörlin, S., Snyder, P. K., Costanza, R., Svedin, U., ... Foley, J. A. (2009). A safe operating space for humanity. *Nature*, 461(7263), 472–475. <https://doi.org/10.1038/461472a>
- Russell, S. (2019). *Human compatible: Artificial intelligence and the problem of control*. Viking.
- Ryan, R. M., & Deci, E. L. (2000). Self-determination theory and the facilitation of intrinsic motivation, social development, and well-being. *American Psychologist*, 55(1), 68–78. <https://doi.org/10.1037/0003-066X.55.1.68>
- Sen, A. (1999). *Development as freedom*. Knopf.
- Sen, A. (2009). *The idea of justice*. Harvard University Press.
- Simon, H. A. (1962). The architecture of complexity. *Proceedings of the American Philosophical Society*, 106(6), 467–482.
- Steffen, W., Richardson, K., Rockström, J., Cornell, S. E., Fetzer, I., Bennett, E. M., Biggs, R., Carpenter, S. R., de Vries, W., de Wit, C. A., Folke, C., Gerten, D., Heinke, J., Mace, G. M., Persson, L. M., Ramanathan, V., Rayers, B., & Sörlin, S. (2015). Planetary boundaries: Guiding human development on a changing planet. *Science*, 347(6223), 1259855. <https://doi.org/10.1126/science.1259855>
- Taleb, N. N. (2012). *Antifragile: Things that gain from disorder*. Random House.
- United Nations. (1966). *International Covenant on Civil and Political Rights*. United Nations Treaty Collection. Retrieved February 19, 2026, from https://treaties.un.org/Pages/ViewDetails.aspx?src=TREATY&mtdsg_no=IV-4&chapter=4&clang=en
- United Nations. (1984). *Convention against Torture and Other Cruel, Inhuman or Degrading Treatment or Punishment*. United Nations Treaty Collection. Retrieved February 19, 2026, from https://treaties.un.org/Pages/ViewDetails.aspx?src=TREATY&mtdsg_no=IV-9&chapter=4&clang=en

United Nations Human Settlements Programme (UN-Habitat). (2003). The challenge of slums: Global report on human settlements 2003. Earthscan. Retrieved February 19, 2026, from <https://unhabitat.org/the-challenge-of-slums-global-report-on-human-settlements-2003>

United Nations Office on Genocide Prevention and the Responsibility to Protect. (2014). Framework of Analysis for Atrocity Crimes. Retrieved February 19, 2026, from <https://www.un.org/en/genocideprevention/publications.shtml>

Üstün, T. B., Kostanjsek, N., Chatterji, S., & Rehm, J. (Eds.). (2010). Measuring health and disability: Manual for WHO Disability Assessment Schedule (WHODAS 2.0). World Health Organization. <https://iris.who.int/handle/10665/43974>

Weber, M. (1978). Economy and society (G. Roth & C. Wittich, Eds.). University of California Press. (Original work published 1922)

World Health Organization. (2001). International classification of functioning, disability and health (ICF). Retrieved February 19, 2026, from <https://www.who.int/standards/classifications/international-classification-of-functioning-disability-and-health>

World Justice Project. (n.d.). Rule of Law Index. Retrieved February 19, 2026, from <https://worldjusticeproject.org/rule-of-law-index/>

APPENDICES

APPENDIX A: SYMBOLS AND NOTATION (v12.0)

Appendix A is normative for symbol meanings and domains. Use the Notation scope rule in Section 0.4 (the specification contract hosted in Appendix AC) when symbols share a glyph; indices, field paths and owning equations distinguish them. Appendix origin stamps record lineage, not expiry of carried-forward provisions.

Terminology note (Normative). *u* indexes Union Scopes (short: Scopes). “Union(s)” remains an allowed legacy alias. This does not change meanings in v9.0 and earlier, it clarifies the canonical term in v9.0.

RealityGroundingRecord / RCR: A PCC-linked record identifying a run’s decision-relevant reality surface, material reality claims, evidence traces, declared unknowns, RealityGroundingStatus, vehicle/content/effect notes where social or institutional constructs are material, surface/territory limitations, and falsification or monitoring triggers. This is the canonical record. The older label RealitySurfaceRecord is deprecated and may be used only as a legacy full-text alias for the surface-description portion of RCR; RSR MUST NOT be used for RealitySurfaceRecord. To avoid acronym collision with ReferenceStructureRecord, use RCR for RealityGroundingRecord and RefStructRecord for ReferenceStructureRecord where abbreviations are needed.

RealityGroundingStatus: The declared status of a decision claim’s contact with the relevant reality surface. Allowed values: GROUNDED, ESTIMATED, ASSUMPTION_BOUND, DECLARED_UNKNOWN, OUT_OF_SCOPE_WITH_RATIONALE, INSUFFICIENT_GROUNDING. Legacy aliases EVIDENCE_BACKED -> GROUNDED and INSUFFICIENT_CONTACT -> INSUFFICIENT_GROUNDING are retained for backward compatibility. At Tier 2 it is operator-declared with review status explicitly recorded; at Tier 3 it must follow the declared reviewer regime; at Tier 4 it must be independently re-derivable from EvidenceTrace and hash-pinned supporting artifacts.

A.1 Sets and Indices

Symbol	Meaning	Domain
<i>u</i>	Union Scope (Scope) index	$u \in \{1,2,3,4,5,6,7\}$
<i>d</i>	Dimension index	$d \in$

Symbol	Meaning	Domain
		$\{1,2,3,4,5,6,7\}$
a	Option (candidate action)	$a \in O$
s	Scenario index	$s \in S$
r	Right index	$r \in R$
k	Impact instance index	integer
g	Protected subgroup index	$g \in G_{\{u,d\}}$

A.2 Core Sets

Set	Definition
U	Operational Scopes {Self, Household, Community, Organization, Polity, Humanity / Global Coordination, Biosphere}
D	Welfare dimensions {Material, Health, Social, Knowledge, Agency, Meaning, Environment}
O	Option set
S	Scenario set
R	Rights set {LIFE, BODY, LBTY, NEED, DIGN, PROC, INFO, ECOL}
$C_r \subseteq U \times D$	Coverage set for right r
$C_{cat} \subseteq U \times D$	Catastrophe cell set for TRC

A.3 Union Stack (Canonical)

Canonical Union Scopes: U1 Self; U2 Household; U3 Community; U4 Organization; U5 Polity; U6 Humanity / Global Coordination; U7 Biosphere. These are the row coordinates of the 7×7 welfare matrix. See Section 2.3 for definitions and Section 4.7 / Appendix AD for cell-level interpretation.

A.3A Scope × Dimension Reference Table (Non-Normative, clarity payload)

The table below reproduces the 7×7 Scope×Dimension reference grid for quick reading and cross-platform stability. The canonical grid definition is presented in Section 4.7.

Scope \ Dimension	D1 Material	D2 Health	D3 Social	D4 Knowledge	D5 Agency	D6 Meaning	D7 Environment
U1 Self	Personal resources, income, essentials	Physical/mental health and safety	Relationships, belonging, support	Learning, understanding, competence	Autonomy, freedom, choice	Purpose, values, spiritual well-being	Local environment affecting self
U2 Household	Household livelihood, housing, food security	Household health, caregiving capacity	Family cohesion, care, safety	Household learning, skills	Household decision rights, roles	Household meaning, culture, identity	Home ecology, consumption footprint

Scope \ Dimension	D1 Material	D2 Health	D3 Social	D4 Knowledge	D5 Agency	D6 Meaning	D7 Environment
U3 Community	Community economy, access to goods/services	Public health, local safety	Social trust, inclusion, cohesion	Education access, knowledge commons	Participation, civic agency	Shared narratives, cultural meaning	Local ecosystems, pollution, resilience
U4 Organization	Organizational resources, viability	Workplace safety, wellbeing	Org culture, fairness, belonging	Org learning, data, competence	Governance, employee agency	Mission coherence, meaning	Operational environmental impacts
U5 Polity	National economy, infrastructure, welfare capacity	Population health, security	Social stability, justice, cohesion	National knowledge, R&D, education	Democratic agency, rights	National identity, legitimacy	National environment, climate policy
U6 Humanity / Global Coordination	Global economy, shared resources	Global health, existential safety	Human solidarity, equity, peace	Global science, open knowledge	Collective governance capacity	Species meaning, long-run purpose	Planetary boundaries, biosphere integrity
U7 Biosphere	Ecosystem services, material substrates	Species health, biodiversity	Symbiosis, interspecies stability	Ecological knowledge, monitoring	Resilience Capacity	Life Continuity	Habitats, climate stability, regeneration

A.4 Welfare Dimensions (Canonical)

Dimension	Name	Summary
D ₁	Material	Resources, infrastructure, subsistence security
D ₂	Health	Physical/mental functioning, morbidity/mortality risk
D ₃	Social	Trust, belonging, relational integrity
D ₄	Knowledge	Epistemic access, learning conditions
D ₅	Agency	Autonomy, effective choice, freedom from coercion
D ₆	Meaning	Purpose/coherence/valued projects (cautious measurement)
D ₇	Environment	Ecological and built context integrity

A.5 Impact Objects

Symbol	Meaning	Range
$\tilde{I}_{dir}(u,d,a)$	Pre-saturation direct impact	$\mathbb{R}$
$I_{dir}(u,d,a)$	Direct impact (post-saturation)	$[-1,+1]$
$\tilde{I}_{prop}(u,d,a)$	Pre-saturation propagated impact	$\mathbb{R}$
$I_{prop}(u,d,a)$	Propagated impact (post-saturation)	$[-1,+1]$
$I_{rights}(u,d,a)$	Worst-off subgroup rights-floor impact	$[-1,+1]$

A.6 Impact Instance Attributes

Symbol	Meaning	Range / Default
μ_k	Magnitude (baseline-delta, normalized)	$[-1, +1]$
r_k	Reach	$[0, 1]$
t_k	Time horizon	$(0, \infty)$ years
ℓ_k	Likelihood	$[0, 1]$
c_k	Confidence	$[0.1, 1]$
e_k	Equity/resilience multiplier	Tier 2: 1.00 only unless named governed instrument; Tier 3: $[0.75, 1.25]$ default 1.00
$h_{\{k,x\}}$	Registered stream coefficient for $x \in \{\text{base, welfare}\}$	$h_{\{k,\text{base}\}}=1$ by definition; $h_{\{k,\text{welfare}\}}=1$ for humans and FPP, otherwise a governed inclusion hypothesis; never use an unregistered s_k or m_k alias
$t_{\min}$	Minimum time horizon floor	0.083 years (≈ 1 month); prevents $\tau(0) = 0$

A.7 Operators and Functions

Symbol

Definition

$(x)^+$

$\max(x, 0)$

$\text{clip}(x, a, b)$

$\max(a, \min(x, b))$

$\text{sat}_\beta(x)$

$\tanh(\beta \times x)$

$\tau(t)$

$\min(1, \ln(1 + \max(t, t_{\min})) / \ln(1 + T_{\text{ref}}))$ Defaults: $T_{\text{ref}} = 25$ years, $t_{\min} = 0.083$ years Cap ensures $\tau(t) \leq 1$ for all t ; floor ensures $\tau > 0$ for all $t > 0$

$\phi(u, d)$

$7(u - 1) + d$ (flattening map for kernel indexing)

$\text{Contain}(u, G_c)$

Decision-specific set of scopes or systems declared to contain, sustain, govern, or enable union scope u in directed containment/dependency graph G_c . Scope order alone does not establish containment.

T_{ref}

Governance reference time horizon for temporal weighting. Default: 25 years. MUST NOT be set < 10 years without charter-level justification.

A.8 Kernel and Propagation

Symbol	Meaning
$K \in$	Ripple kernel matrix (sparse)

Symbol	Meaning
$\mathbb{R}^{49 \times 49}$	
$K_{\{ij\}}$	Kernel entry mapping source j to target i (target-row, source-column)
$k_{\max}$	Maximum absolute kernel entry bound, default 0.5
$k_{\text{use_min}}$	Minimum absolute kernel edge magnitude for “relied-upon edge” classification in KQS computation (Section 6.5.1). Default: 0.05 unless the PCC declares a stricter value.
$\rho_{\max}$	Maximum absolute row-sum bound, default 0.9
$\ K\ _{\infty}$	Maximum absolute row-sum norm of K : $\ K\ _{\infty} := \max_i \sum_j K_{\{ij\}} $. Used for Tier 1-3 kernel stability verification (Section 6.4).
$\rho(K)$	Spectral radius of K . For Tier 1-3, treated as satisfied when $\ K\ _{\infty} < 1$ per Section 6.4 (since $\rho(K) \leq \ K\ _{\infty}$).
NONE mode	$I_{\text{prop}} := I_{\text{dir}}$
QUICK mode	$\tilde{I}_{\text{prop}} = I_{\text{dir}} + K \times I_{\text{dir}}$, then saturate

Stream note (Normative). Kernel propagation applies per stream $x \in \{\text{base}, \text{welfare}\}$ as specified in Appendix B. In Appendix A, the symbols I_{dir} , $\tilde{I}_{\text{prop}}$, and I_{prop} may be used as shorthand for their stream-specific forms $I_{\text{dir},x}$, $\tilde{I}_{\text{prop},x}$, and $I_{\text{prop},x}$. Admissibility computation uses the Base stream ($h_k:=1$); RLS uses the Welfare stream with the governed welfare-inclusion hypothesis set carried through robust sensitivity analysis.

A.9 Rights and Constraints

Symbol	Meaning
θ_r	Rights threshold for right r
$v_r(a)$	Violation depth for right r under option a
$\text{NCRC}(a)$	Rights admissibility predicate
$\text{TRC}(a)$	Tail-risk admissibility predicate
τ_{TRC}	TRC corridor threshold
α	CVaR tail level $\in (0, 1)$
$\text{CVaR}_{\alpha}[L(a)]$	Conditional Value-at-Risk for loss distribution of option a
θ_{pos}	Positive impact threshold (containment trigger), default 0.05
τ_c	Containment tolerance threshold, default -0.10
D_c	Containment ancestor depth, default 2
$I_{\text{RF}}(u,d,a g,r)$	Subgroup-specific rights-floor impact under Section 7.3; B.8 takes its minimum over the declared G_{RF} set. It is not a positive welfare average.
b_{tail}	Tail probability mass $1 - \alpha$ in Appendix D.7; distinct from the direct-saturation parameter β .

A.10 Scoring and Uncertainty

Symbol	Meaning
w_u	Union weight ($w_u \geq 0$, $\sum_u w_u = 1$)
v_d	Dimension weight ($v_d \geq 0$, $\sum_d v_d = 1$)
$m(u,d)$	Applicability mask for RLS only $\in \{0, 1\}$
$\kappa(u,d)$	Cell multiplier for RLS, default 1
$RLS(a)$	RippleLogic Score
$\sigma(u,d,a)$	Cell-level uncertainty proxy
$\sigma_{RLS}(a)$	RLS uncertainty proxy
δ	Discrimination threshold for decisive lead, default 2
ϵ	Stabilizer in Gap computation, default 10^{-6}
λ	Risk-aversion coefficient for RLS_{adj} , default 0.5

A.11 Structural Metrics

Symbol	Meaning
UCI_u	Union Coherence Index for union $u \in [0, 1]$
$\Delta UCI_u(a)$	Coherence change under option $a = UCI_u(a) - UCI_u(\text{baseline})$
HOI_t	Hollowing-Out Index at time t
H_u	Cohesion component of UCI
F_u	Flow component of UCI
R_u	Resilience component of UCI
E_u	Equity component of UCI; NA where genuinely inapplicable, UNKNOWN where applicable evidence is missing
$X_{\{u,j\}}$	Applicable structural component value for scope u in B.15; j identifies H, F, R or E under the declared instrument.
s_{RLS}, s_{UCI}	Positive preregistered reference scales for HOI in B.17; missing or unstable scales make that HOI claim unavailable.

A.12 Audit Artifacts

Symbol	Meaning
PCC	Provenance and Compliance Certificate
5SPR	Five-Sentence Public Rationale
AIL	Artifact Integrity Law (Tier 2-3 integrity rules)
Audit	Standard labels for invalidity or warnings (Appendix H)

Symbol	Meaning
flags	

A.13 Scenario and TRC Objects

Symbol	Meaning
S	Scenario set
p_s	Scenario probability ($p_s \geq 0$, $\sum_s p_s = 1$)
$L(a,s)$	Scenario loss for option a under scenario s
ω_c	Catastrophe cell weight ($\omega_c \geq 0$, $\sum_c \omega_c = 1$)
p_{floor}	Minimum probability floor per mandatory tail category, default 0.02

A.14 SGP v8.8 Integration

Symbol	Meaning
MPS(E,t)	Moral Patienthood Surface: evidence families, band, interval, confidence, and protection posture.
RMCP(C,D,t)	Reality-Management Capacity Profile for a declared entity/collective C, domain D, and time t; informative only.
P100(C)	Open top-plateau status for functional equivalence to or greater integrated reality-management capacity than the human collective anchor.
FPP(E)	Full Protection Plateau: constitutional or governed non-downgrade rights-of-protection status.
GPR(E,r,t)	Governance Participation Readiness for a named process or institution.
SPR(E,r,t)	Stewardship/Power Readiness verdict for a named consequential role.
ICP(E,D,t)	Optional Intelligence Capacity Profile for a declared domain; informative only.
h_k	Welfare-inclusion hypothesis coefficient used in robust RLS sensitivity; ordinarily 0 or 1, with $h_k=1$ for humans and FPP entities.

End Appendix A.

APPENDIX B: CANONICAL EQUATIONS (TIER 1-3 EXECUTABLE)

Appendix B is normative and is the canonical equation pack for Tier 1-3 implementations. All equations are reproduced explicitly for standalone executability.

B.1 Temporal Weighting

Defaults: $T_{\text{ref}} = 25$ years; $t_{\text{min}} = 0.083$ years (≈ 1 month).

$$\tau(t) = \min(1, \ln(1 + \max(t, t_{\text{min}})) / \ln(1 + T_{\text{ref}}))$$

This function has three properties by construction:

(a) Logarithmic growth ensures that short-term effects receive proportionally less weight than long-term effects, without the exponential discounting that aggressively devalues future generations.

(b) The $\min(1, \cdot)$ cap ensures that no effect receives more temporal weight than the governance reference horizon. At T_{ref} , $\tau = 1.00$. Beyond T_{ref} , τ remains 1.00 (capped).

(c) The $\max(t, t_{\min})$ floor ensures that $\tau(t) > 0$ for all positive t , preventing very short-term effects from being zeroed out.

Illustrative values (with defaults $T_{\text{ref}} = 25$, $t_{\min} = 0.083$):

At 100 years: $\tau = 1.00$ (capped)

At 50 years: $\tau = 1.00$ (capped)

At 25 years: $\tau = 1.00$

At 10 years: $\tau \approx 0.74$

At 5 years: $\tau \approx 0.55$

At 1 year: $\tau \approx 0.21$

At $t_{\min}$ (≈ 1 month): $\tau \approx 0.025$

Historical change note (v7.5.0; carried forward unchanged in v9.0): This replaces the uncapped function $\tau(t) = \ln(1+t)/\ln(1+T_{\text{ref}})$ from v7.4.5, which produced $\tau(50) \approx 1.21$. The capped version is normatively preferred. See Section 5.3 for design rationale.

B.2 Direct Impact Aggregation (Pre-Saturation; stream-separated)

B.2A Base stream (admissibility; $h_k:=1$):

$$\tilde{I}_{\text{dir_base}}(u,d,a) = \sum_{\{k \in K(u,d,a)\}} [r_k \times \tau(t_k) \times \ell_k \times c_k \times e_k \times (1) \times \mu_k]$$

B.2A Gate-critical confidence guard. In the Base stream used by RF/NCRC, TRC, and CSV, c_k represents evidence support but cannot by itself excuse unresolved severe adverse gate-critical impacts. Apply $\text{GateAdverseBound}(k)$ from Section 5.2: use the most adverse admissible harm endpoint and least favourable gate-supporting endpoint; if only c_k is available, set $c_{\text{gate},k}:=1$ for adverse gate computation and retain $c_{\text{gate},k}:=c_k$ for beneficial safety claims. If the bound is not reproducible and the gate could change, return UNKNOWN, ESCALATE, NARROW, or REFUSE rather than PASS.

B.2B Welfare stream (RLS; welfare-inclusion hypotheses):

$$\tilde{I}_{\text{dir_welfare}}(u,d,a;h) = \sum_{\{k \in K(u,d,a)\}} [r_k \times \tau(t_k) \times \ell_k \times c_k \times e_k \times h_{\{k,\text{welfare}\}} \times \mu_k]$$

Here $h_{\{k,\text{welfare}\}}^{\{j\}}$ is selected from the governed welfare-inclusion hypothesis set only for a declared sensitivity evaluation j ; the Base stream (B.2A) fixes $h_k:=1.0$ for all instances. MPS bands or intervals MUST NOT be substituted directly for h_k . A representative display may be shown only if all decision-material hypothesis results are also reported.

Where $K(u,d,a)$ is the set of impact instances asserted for cell (u,d) under option a .

B.3 Direct Saturation (per stream)

For $x \in \{\text{base}, \text{welfare}\}$:

$$I_{\text{dir}_x}(u,d,a) = \tanh(\beta \times \tilde{I}_{\text{dir}_x}(u,d,a))$$

Default: $\beta = 2$.

B.4 Missing Data (Ignorance Penalty Phantom Instance)

If cell (u,d) is required-active but $K(u,d,a) = \emptyset$, add phantom instance with:

Parameter	Phantom Value
μ_{phantom}	-0.10

Parameter	Phantom Value
r	1
t	T_ref (25 years)
ℓ	1
c	1
e	1
h_{k,base}	1
h_{k,welfare}	1

B.5 Flattening Map (Kernel Indexing)

Vectorize I_{dir} into a 49-vector by:

$$i = \phi(u,d) = 7(u - 1) + d$$

B.6 Propagation Modes (Tier 1-3: NONE and QUICK only)

NONE:

For $x \in \{\text{base}, \text{welfare}\}$: $I_{prop_x} := I_{dir_x}$

QUICK:

For stream $x \in \{\text{base}, \text{welfare}\}$: $\tilde{I}_{prop_x} = I_{dir_x} + K \times I_{dir_x}$

Then apply post-propagation saturation (B.7).

B.7 Post-Propagation Saturation (per stream)

For $x \in \{\text{base}, \text{welfare}\}$:

$$I_{prop_x}(u,d,a) = \tanh(\beta_{prop} \times \tilde{I}_{prop_x}(u,d,a))$$

Default: $\beta_{prop} = 1$.

Applicability rule. B.7 applies only when a propagation kernel is used (QUICK or a later governed kernel). Under NONE, $I_{prop_x} := I_{dir_x}$; no second post-propagation saturation is applied.

B.8 Worst-Off Subgroup Impact (Rights)

$$I_{rights}(u,d,a,r) = \min_{\{g \in G_{RF}(u,d,a,r)\}} I_{RF}(u,d,a|g,r)$$

B.9 NCRC Violation Depth and Rights Admissibility

For each right r , define three independent violation channels using Section 7.4 inputs:

The severe-hazard activation predicate, $h_{sev,k,r}$, $\eta_{r,k}$, event-class screening and scoped tolerance $\rho_{\{k,r\}}$ follow Section 7.4.3. A shorthand ρ_r does not replace the event-, population- and authority-specific tolerance record. Missing activation or tolerance evidence remains unresolved under the reviewed-empty-set convention.

$$v_r^{\wedge floor(a)} = \max_{\{u,d\} \in C_r} (\theta_r - I_{rights}(u,d,a))^{\wedge +}$$

$$v_r^{\wedge cat(a)} = \max_{\{k \in K_r(a)\}} z_{\{k,r\}}$$

$$v_r^{\text{risk}}(a) = \max_{\{k \text{ in } K_r(a): \text{SevereRightsHazardActive}(k,r)\}} (p_{\{k,r\}}^U - \rho_{\{k,r\}})^{++}$$

where $(x)^{++} = \max(x, 0)$, $z_{\{k,r\}}$ is the categorical-prohibition indicator, $h_{\{\text{sev},k,r\}}$ is intrinsic rights severity, $p_{\{k,r\}}^U$ is the conservative upper event-probability bound, and ρ_r is the governed rights-risk tolerance. The risk-check index includes both individual constraints and required aggregate event-family probability records under Section 7.4.4. Each uses its scoped $\rho_{\{k,r\}}$; aggregation does not add a welfare impact or dilute a categorical prohibition.

If any required severe-hazard field is missing, the right receives NCRC_UNKNOWN; missing data MUST NOT be converted to zero.

NCRC(a) = TRUE if and only if, for every protected right r, $v_r^{\text{floor}}(a)=0$, $v_r^{\text{cat}}(a)=0$, $v_r^{\text{risk}}(a)=0$, and no material rights field is unresolved.

For Emergency Mode comparison, use the ordered compound tuple $q_r(a) = (v_r^{\text{cat}}(a), v_r^{\text{risk}}(a), v_r^{\text{floor}}(a))$ within the canonical right-priority order. A lower welfare score or shorter duration cannot rescue a non-zero categorical, severe-risk, or floor component.

B.10 TRC Loss (Tier 1-3 Bounded-Impact Mode)

Let C_{cat} be catastrophe cells and ω catastrophe weights.

Normative stream binding (Design A): TRC MUST be computed from the Base stream ($I_{\text{prop_base}}$). TRC MUST NOT use Welfare-stream (sentience-weighted) impacts.

For scenario s:

$$L(a,s) = \sum_{\{c \in C_{\text{cat}}\}} \omega_c \times (-I_{\text{prop_base},c}(a|s))^{+}$$

TRC uses CVaR (Appendix D for discrete computation) and corridor threshold τ_{TRC} :

TRC(a) = TRUE if and only if $\text{CVaR}_{\alpha}[L(a)] \leq \tau_{\text{TRC}}$, every active class-specific cap or tolerance under Section 8.4A is satisfied, no unresolved material C3 or C4 field remains, and all mandatory scenario-family requirements are satisfied or governed exceptions are recorded.

B.11 Containment Mode A (Binding)

Containment Mode A is one required structural subtest inside the broader CSV resolver; it is not the whole CSV Gate.

Define positively moving union scopes:

$$S_u(a) = \sum_d v_d^{\text{cont}} * I_{\text{prop_base}}(u,d,a)$$

Containment trigger weights use the containment-specific vector v_d^{cont} (default uniform 1/7). HDW welfare weights MUST NOT be used for containment trigger detection.

$$U_{\text{pos}}(a) = \{u : S_u(a) \geq \theta_{\text{pos}}\}$$

Let G_c be the decision-specific directed containment/dependency graph. Increasing union-scope order is not sufficient evidence of containment. For each u in $U_{\text{pos}}(a)$, let $\text{Contain}(u, G_c)$ be the set of declared scopes or systems whose integrity materially contains, sustains, governs, or enables the benefit. Where UCI is an applicable instrument:

$$M_u(a) = \min_{\{u' \text{ in } \text{Contain}(u, G_c)\}} \Delta \text{UCI}_{\{u'\}}(a)$$

If $\text{Contain}(u, G_c)$ is empty for a triggered scope, the containing-system map is incomplete. The run MUST record CONTAINMENT_MAP_INCOMPLETE and escalate or leave the CSV subcondition unresolved; the minimum over an empty set MUST NOT pass vacuously.

and the UCI subpredicate passes only when $M_u(a) \geq \tau_c$ for every triggered scope.

The complete CSV status is resolved from all material subconditions:

$CSV_status(a) = \text{Resolve}(\text{containment_integrity}, \text{structural_viability}, \text{source_coupling}, \text{physical_causal_admissibility}, \text{dependency_and_lock_in}, \text{reversibility}, \text{monitoring}, \text{controls}, \text{authority_and_legitimacy}, \text{evidence_sufficiency})$

CSV_PASS or CSV_PASS_WITH_CONTROLS requires every material subcondition to be PASS, PASS_WITH_CONTROLS, or a documented NOT_MATERIAL disposition. Any material FAIL, unresolved gate-critical UNKNOWN, missing required profile, or structurally non-viable condition prevents ordinary selectability.

CSV_NOT_MATERIAL remains a documented pass-equivalent disposition only when the PCC demonstrates that no CSV subcondition is material to the declared claim boundary. UCI is therefore one possible required subpredicate, not an if-and-only-if definition of CSV.

B.12 RippleLogic Score (RLS)

$$q(u,d) = w_u \times v_d \times m(u,d) \times \kappa(u,d)$$

$$Q = \sum_u \sum_d q(u,d)$$

$$RLS(a) = [\sum_u \sum_d q(u,d) \times I_{\text{prop_welfare}}(u,d,a)] / Q, \text{ for } Q > 0$$

If $Q = 0$, RLS is undefined. Return RLS_NO_ACTIVE_MASS, repair the mask/weight declaration, narrow the claim, or refuse deterministic ranking; zero MUST NOT be imputed.

Default: $\kappa(u,d) = 1$, $m(u,d) = 1$ unless masked.

B.12A κ governance note (Normative). Non-default $\kappa(u,d)$ values are allowed only within the bounded governance interval declared in Section 10.1 and MUST preserve sign, rights-floor protection, catastrophe visibility, and replayability. Validator checks SHALL reject any $\kappa(u,d)$ that acts as a hidden mask, hidden union-floor rewrite, or undisclosed prominence override.

B.13 RLS Uncertainty (Optional; Tier 3 Required)

Definition of $\sigma(u,d,a)$ (normative minimum): PCC MUST declare one of the following cell-level uncertainty proxies and apply it consistently across options:

When Method B is used for a unique-selection claim, the transformation-consistent uncertainty requirement in Section 10.3.2 controls. A reproducible legacy support-proxy calculation is not by itself a calibrated final-impact uncertainty warrant.

Method A (interval half-width): If the PCC records a confidence interval $[L,U]$ for the cell impact $I(u,d,a)$, set $\sigma(u,d,a) = (U - L)/2$.

Method B (confidence-derived contribution-mass rule): If the PCC records impact instances with confidences $c_k \in [0.1,1]$ and selects Method B, derive $c(u,d,a)$ using Section 10.3.2. Let $A_{\text{cell}}(u,d,a) = \min(1, \sum_k |x_k|)$, where x_k is the pre-confidence contribution after reach, temporal, likelihood, adjustment, and permitted sentence terms and before saturation. Then: $\sigma(u,d,a) = (1 - c(u,d,a)) \cdot A_{\text{cell}}(u,d,a)$

Phantom uncertainty (Normative). For a governed phantom under Section 5.7, the confidence used in the point estimate remains 1, but the Method B uncertainty confidence is 0; hence its support proxy is $\sigma = A_{\text{cell}}$. This does not turn a missing observation into an observed zero or waive the Section 10.3.2 transformation-consistency and dependence requirements for a unique-selection claim.

Method C (calibrated table): Use a pre-registered mapping from qualitative uncertainty labels (for example, LOW/MED/HIGH) to σ values, stored in the ProofPack or PCC appendix.

σ values MUST be clipped to [0,1]. Unless a stricter governed method is declared, $\sigma(u,d,a)$ is interpreted as the pre- κ cell-level uncertainty input and SHALL be aggregated into σ_{RLS} using the same w_u , v_d , $m(u,d)$, and $\kappa(u,d)$ posture used for RLS.

$$\sigma_{\text{RLS}}(a) = \sqrt{[\sum_u \sum_d (q(u,d) \times \sigma(u,d,a))^2] / Q}, \text{ for } Q > 0$$

This is the independent-cell approximation. Tier 3 and decision-material Tier 2 runs MUST declare dependence clusters and compute the perfect-correlation or covariance-aware stress specified by the Dependence-cluster uncertainty rule in Section 10.3. If ranking or decisiveness changes, set `RLS_DEPENDENCE_SENSITIVE` and treat the result as non-decisive until governed resolution. The declared cell uncertainty MUST also include material measurement or elicitation resolution; exact zero is allowed only for exact deterministic inputs.

Risk-adjusted score (optional):

$$\text{RLS}_{\text{adj}}(a) = \text{RLS}(a) - \lambda \times \sigma_{\text{RLS}}(a)$$

By default, `RLS_adj` is diagnostic only and MUST NOT replace RLS for ranking or decisiveness unless the run declares ex ante that risk-adjusted ranking is the governing ranking basis. If such a declaration is made, the same declared basis MUST be used consistently for ranking, Gap, and decisiveness across all compared options in the run.

B.14 Discrimination Gap (Decisive vs Non-Decisive)

$$\text{Gap}(a,b) = \text{abs}(\text{RLS}(a) - \text{RLS}(b)) / \sqrt{\text{sigma}_{\text{RLS}}(a)^2 + \text{sigma}_{\text{RLS}}(b)^2 + \text{epsilon}}$$

$\text{SignedGap}(a,b) = (\text{Score}(a) - \text{Score}(b)) / \sqrt{\text{sigma}_{\text{RLS}}(a)^2 + \text{sigma}_{\text{RLS}}(b)^2 + \text{epsilon}}$. Score is the ex ante governing basis in Section 10.4: ordinary RLS by default, or the expressly declared risk-adjusted basis. The absolute Gap reports separation magnitude; only the signed leader-versus-every-rival comparison establishes the required direction.

Default: $\text{delta} = 2$, $\text{epsilon} = 1\text{e-}6$.

A lead is decisive only when the Section 10.4 signed leader comparison exceeds delta against every other selectable option under every required variant, the same leader survives, and the Section 10.4B robustness record is complete. Any failed comparison or material sensitivity blocks a unique decisive claim.

Demonstration-only or unvalidated uncertainty cannot support `ALLOW_FRAMEWORK_SELECTION` or `SELECTED_DECISIVE`. Otherwise retain unresolved contenders and `REFUSE_DETERMINISTIC_SELECTION`; any governed tie-break preference and separate `AuthoritySelectionRecord` follow Section 11.5. A sole survivor has no pairwise Gap.

B.14A Selection-Integrity Equations and Predicates (Normative)

Event-probability residual welfare: for scenario-conditioned outcome $X_{ud}(a,s)$, $I_{ud}(a) = \sum_s p_s \cdot \tanh(\beta \cdot X_{ud}(a,s))$. Expected-frequency/exposure and already-expected effects follow Section 5.2B.

Scenario-order binding. Under NONE, the scenario-specific direct saturated value is also its propagated value. Under QUICK, construct and saturate each scenario's direct Base/Welfare field, apply B.6 and B.7 within that scenario, and then probability-average the final Welfare impacts for RLS. TRC retains each scenario's loss distribution for CVaR rather than collapsing it into that welfare average. Do not exchange nonlinear saturation with averaging.

Rights non-dilution: $v_{\text{floor}}(a) = \max(v_{\text{floor}}^{\text{cell}}(a), \max_{j,r} v_{j,r}^{\text{token}}(a))$. NCRC passes only when floor, categorical, and severe-hazard channels are all zero and no required field is unresolved.

Comparison-mask symmetry: $m_run(u,d)$ is common across compared options. An unjustified option-specific mask makes comparative RLS invalid.

Robust selection predicate: $UniqueSelect(a^*)=TRUE$ only if a^* is decisively separated from every other selectable option under Section 10.4 in every completed required variant, every registered module is `PASS_NO_REVERSAL` or `NOT_TRIGGERED_WITH_RATIONALE`, and no upstream gate or authority rule blocks selection.

Normative underdetermination predicate: if the evidence surface is sufficiently closed and two legitimate ex-ante weight constitutions select different options, return `NORMATIVE_UNDERDETERMINED` rather than expanding σ_RLS .

B.15 UCI Computation

Let A_u be the set of UCI components applicable to Union Scope u . Then:

$$UCI_u = [\sum_{\{j \in A_u\}} \alpha_j X_{\{u,j\}}] / [\sum_{\{j \in A_u\}} \alpha_j]$$

Default: $\alpha_H = \alpha_F = \alpha_R = \alpha_E = 0.25$ before applicable-component renormalization. NA components are excluded from A_u ; UNKNOWN components are not scored and make UCI unavailable for the affected claim. For Self (U1), Equity is ordinarily NA unless a governed within-self equity instrument is material. For Biosphere (U7), Equity is ordinarily NA unless a governed ecological-distribution or ecological-justice instrument exists. No neutral or perfect value may be inserted for NA or UNKNOWN. If A_u is empty, UCI_u is UNAVAILABLE; the denominator is not evaluated and no value may be imputed.

B.16 ΔUCI Computation

$$\Delta UCI_u(a) = UCI_u(a) - UCI_u(\text{baseline})$$

B.17 HOI Computation

$$HOI_t = EMA_{\lambda}(\Delta RLS_t / s_RLS)_t - EMA_{\lambda}(\Delta UCI_t / s_UCI)_t$$

Scale rule (Normative). s_RLS and s_UCI MUST be positive, preregistered reference scales derived from a declared reference range or robust dispersion estimate and held fixed for the review window. If either scale is unavailable or unstable, HOI is UNAVAILABLE; raw unstandardized changes MUST NOT be subtracted.

Default half-life: 3 periods.

End Appendix B.

APPENDIX C: RIGHTS / NCRC CANON PACK (TIER 1-3)

Appendix C is normative.

C.1 Rights Thresholds (Canonical)

Right	Code	Threshold θ_r
Life	LIFE	-0.90
Bodily Integrity	BODY	-0.70
Liberty	LBTY	-0.65
Basic Needs	NEED	-0.50
Dignity	DIGN	-0.55
Due Process	PROC	-0.45

Right	Code	Threshold θ_r
Information	INFO	-0.40
Ecological Integrity	ECOL	-0.65

Normative interpretation rule: Thresholds θ_r are admissibility floors, not weights. They do not represent “how important” a right is; they represent the minimum allowable protection level.

Threshold-scale clarification (Normative). θ_r is defined on the post-saturation rights-impact scale I_{rights} . For an unsaturated single-stream comparison, the corresponding pre-saturation aggregate is $\text{atanh}(\theta_r) / \beta_{RF}$. With the canonical $\beta_{RF} = 2$, $\tanh(\beta_{RF} x)$ bounds values to $(-1, +1)$ but is expansive in magnitude near zero and compressive only near the bounds; it is not a small-signal attenuator. Threshold calibration and sensitivity review MUST state the scale used.

C.2 Rights Coverage Sets C_r (Authoritative)

Right	Coverage Set C_r
LIFE	$\{(u, \text{Health}): u \in \{1,2,3,4,5,6\}\} \cup \{(6, \text{Environment})\}$
BODY	$\{(u, \text{Health}): u \in \{1,2,3,4,5,6\}\}$
LBTY	$\{(u, \text{Agency}): u \in \{1,2,3,4,5,6\}\} \cup \{(u, \text{Social}): u \in \{3,4,5,6\}\}$
NEED	$\{(u, \text{Material}): u \in \{1,2,3,4,5,6\}\} \cup \{(u, \text{Health}): u \in \{1,2,3,4,5,6\}\}$
DIGN	$\{(u, \text{Social}): u \in \{1,2,3,4,5,6\}\} \cup \{(u, \text{Agency}): u \in \{1,2,3,4,5,6\}\}$
PROC	$\{(u, \text{Agency}): u \in \{4,5,6\}\} \cup \{(u, \text{Knowledge}): u \in \{4,5,6\}\} \cup \{(u, \text{Social}): u \in \{4,5,6\}\}$
INFO	$\{(u, \text{Knowledge}): u \in \{1,2,3,4,5,6\}\} \cup \{(u, \text{Agency}): u \in \{1,2,3,4,5,6\}\}$
ECOL	$\{(6, \text{Environment}), (7, \text{Environment})\}$

Rights-routing clarification for D6 and D7 (Normative). Appendix C.2 defines the cells directly covered by each canonical right. A D6 Meaning or D7 Environment entry does not acquire rights-floor status merely because it is rights-relevant. When the underlying effect instantiates a protected interest — for example freedom of thought, conscience, religion, or nonbelief; coercion or discrimination; dignity; information access; health; basic needs or habitability; due process; or protected life-support conditions — the run MUST construct and evidence the corresponding rights-floor instance in the applicable canonical covered dimension(s), or declare a governed rights-coverage extension under Section 12.8. It MUST NOT leave a materially protected effect only as compensable D6/D7 welfare. Ordinary purpose, identity, place attachment, environmental amenity, or ecological condition is not automatically a rights violation.

C.3 Emergency Mode Rights Priority (Canonical)

When Emergency Mode is invoked ($A_{NCRC} = \emptyset$), rights are ordered lexicographically:

[LIFE, BODY, ECOL, LBTY, NEED, DIGN, PROC, INFO]

This ordering MUST be used exactly for lexicographic minimization of the violation depth vector.

C.4 NCRC Violation Depth Definition (Canonical)

For every right r , record the compound violation tuple:

$$q_r(a) = (v_r^{\text{cat}}(a), v_r^{\text{risk}}(a), v_r^{\text{floor}}(a))$$

with:

- $v_r^{\text{floor}}(a) = \max_{\{(u,d) \in C_r\}} (\theta_r - I_{\text{rights}}(u,d,a))^+;$
- $v_r^{\text{cat}}(a) = \max_{\{k \in K_r(a)\}} z_{\{k,r\}};$
- $v_r^{\text{risk}}(a) = \max_{\{k \in K_r(a) : \text{SevereRightsHazardActive}(k,r)\}} (p_{\{k,r\}}^U - \rho_{\{k,r\}})^+.$

A missing required severe-hazard probability, tolerance, categorical-status field, or gate-critical evidence disposition produces NCRC_UNKNOWN, not zero. Apply the scoped tolerance and aggregate event-family probability requirements of Section 7.4.4 before the severe-risk maximum; unresolved material dependence or exposure boundaries cannot receive an implied pass.

C.5 NCRC Predicate (Canonical)

NCRC(a) = TRUE if and only if, for every right r, all three components of $q_r(a)$ equal zero and no material rights evidence field is unresolved.

The NCRC passing set is $A_{\text{NCRC}} = \{a \in O : \text{NCRC}(a) = \text{TRUE}\}$. Any categorical prohibition, severe-hazard exceedance, floor violation, or required unresolved field blocks an ordinary NCRC pass.

C.6 Subgroup Conservative Bound (Canonical)

When subgroup disaggregation is infeasible for a rights-covered cell:

Let $\gamma_{\text{subgroup}} = 1.5$ (default conservatism factor).

- If $I_{\text{RF_base}}(u,d,a,r) < 0$: set $I_{\text{rights}}(u,d,a,r) = \max(-1, \gamma_{\text{subgroup}} \times I_{\text{RF_base}}(u,d,a,r))$
- Else: set $I_{\text{rights}}(u,d,a,r) = I_{\text{RF_base}}(u,d,a,r)$

This bound applies only for NCRC checking, not for RLS scoring.

Emergency comparison tie rule (Normative). The saturated I_{rights} value remains controlling for ordinary admissibility. If two or more Emergency Mode options tie on v_r^{floor} because the conservative clamp maps them to -1 , the run MUST use the unclamped conservative value $\gamma_{\text{subgroup}} \times I_{\text{RF_base}}$ as a secondary ordering key for that right, with more negative values treated as worse and the trace recorded in the PCC. This secondary key does not convert any failing option into an ordinary NCRC pass.

If $I_{\text{RF_base}}$ cannot be constructed from reviewed rights-floor instances, the cell is NCRC_UNKNOWN; the run MUST collect evidence, narrow, escalate, or refuse rather than reuse a temporally or confidence-attenuated welfare value as a rights pass.

C.7 Emergency Mode Documentation Requirements (Canonical)

When Emergency Mode is invoked, the PCC MUST include:

- Emergency declaration and trigger conditions;
- the complete $q_r(a) = (v_r^{\text{cat}}, v_r^{\text{risk}}, v_r^{\text{floor}})$ tuple for each option and each right;
- the lexicographic comparison trace using the Section C.3 right-priority order and component order categorical $\rightarrow$ severe risk $\rightarrow$ floor depth;
- all NCRC_UNKNOWN fields and the conservative disposition used;
- independent challenger attestation, or CHALLENGE_DEFERRED_EMERGENCY;
- mitigation/remediation plan with timeline;
- monitoring, appeal, remedy, exit, review cadence, and return-to-normal triggers.

Emergency Mode does not convert a rights-violating option into an ordinary NCRC pass. It documents a governed provisional response under necessity.

End Appendix C.

APPENDIX D: TRC AND SCENARIO GOVERNANCE CANON PACK (TIER 1-3)

Appendix D is normative.

D.1 Base Catastrophe Cell Set

$C_cat_base = \{(6, \text{Health}), (6, \text{Environment}), (7, \text{Environment})\}$

These correspond to:

- (6, Health): Humanity / Global Coordination-Health, representing global-scale health viability
- (6, Environment): Humanity / Global Coordination-Environment, representing civilization-scale habitability conditions
- (7, Environment): Biosphere-Environment, representing Earth-system integrity

D.1A Default Governance-Lock-In Extension Profile (Normative)

Trigger condition: This profile **MUST** be activated when any of the following is plausibly in scope for the decision: surveillance normalization; emergency-power persistence beyond justified sunset; due-process suspension or durable recourse collapse; information-ecosystem capture; irreversible dependency lock-in on critical infrastructure or governance functions; or concentrated arbitrary control without meaningful exit, review, or contestability.

Default governance-lock catastrophe cell set:

$C_cat_lock_default = \{(5, \text{Agency}), (5, \text{Knowledge}), (6, \text{Agency}), (6, \text{Knowledge})\}$

Default weighting rule when active:

If no stricter governed profile is declared:

Evaluate C_cat_base and $C_cat_lock_default$ as separate active catastrophe profiles. Normalize weights within each profile, compute profile-specific loss and CVaR, and require every active profile to satisfy its governed corridor. A combined union summary may be reported only as non-controlling information and cannot rescue a failed profile.

Interpretation:

(5, Agency): polity-scale loss of contestability, due process, or civic agency

(5, Knowledge): polity-scale degradation of information integrity and epistemic access

(6, Agency): civilization-scale agency loss, lock-in, or durable corrigibility collapse

(6, Knowledge): civilization-scale information capture, knowledge bottlenecking, or durable epistemic enclosure

Symmetry rule: The same active catastrophe set and weights **MUST** be applied across all compared options in the run.

PCC rule: The PCC **MUST** record whether the governance-lock profile was inactive, active by default trigger, or replaced by a stricter governed profile.

D.2 Mandatory Tail Scenario Categories (Canonical)

The scenario set **MUST** include the following tail categories unless explicitly justified as not plausible.

Governance / lock-in catastrophe is a mandatory tail category if and only if the governance-lock-in extension profile in D.1A is active; otherwise it is recommended but not mandatory. When that profile is active, the scenario set **MUST** include at least one governance-lock-in tail category with category probability floor $p_floor \geq 0.02$.

- Pandemic/biological disruption: Disease outbreak with significant mortality/morbidity
- Climate tipping cascade: Triggering of climate tipping points with cascading effects
- Financial system collapse: Systemic failure of financial infrastructure
- Major conflict escalation: Armed conflict with regional or global implications
- Critical infrastructure failure: Failure of essential services (power, water, communications)

Governance / lock-in catastrophe: Surveillance normalization, durable recourse collapse, information-ecosystem capture, irreversible dependency lock-in, or concentrated arbitrary control over critical governance or infrastructure functions

D.3 Probability Floor (Canonical)

Minimum category probability floor: $p_{\text{floor}} \geq 0.02$

Meaning: The sum of probabilities of scenarios in each mandatory category must be ≥ 0.02 unless explicitly justified as implausible for the specific decision context per Section 8.3.2.

D.4 Scenario Set Size Requirements (Canonical)

Tier	TRC Required?	Minimum
Tier 1	Required qualitative trigger screen	N/A
Tier 2	Required when catastrophe relevance plausible	≥ 5 minimum
Tier 3	Required	≥ 20 minimum
Tier 4	Design target only	N/A

D.5 Scenario Definition Template (Normative Minimum Fields)

Each scenario record in PCC MUST include:

Field	Description
Scenario ID	Unique identifier within the run
Name	Descriptive name
Category	One of the mandatory categories or "baseline/other"
Narrative	2-5 sentences describing the scenario
Time horizon	Timing assumptions
Key stressors	What breaks, how, and cascading effects
Parameterization	What changes in impacts/likelihoods under this scenario
Probability p_s	Assigned probability with provenance
Provenance	Data source, model, expert elicitation, or governance prior
Severity class	C1_SEVERE / C2_SYSTEMIC / C3_CIVILIZATION_THREATENING / C4_EXISTENTIAL
Dependence cluster	Declared common-cause or dependence-cluster identifier; NONE only with justification.

D.6 Loss Construction (Bounded-Impact Mode, Tier 1-3 Canon)

Normative stream binding (Design A): TRC MUST be computed from the Base stream (I_{prop_base}). TRC MUST NOT use Welfare-stream (sentence-weighted) impacts.

For each scenario s :

$$L(a,s) = \sum_{c \in C_cat} \omega_c \times (-I_{prop_base,c}(a|s))^+$$

with $\sum_{c \in C_cat} \omega_c = 1$.

Default catastrophe weights: Uniform over C_cat unless otherwise governed.

Any non-default ω_c vector MUST be declared ex ante, normalized within the active profile, justified, applied symmetrically, and accompanied by the uniform-weight counterfactual when the run is Tier 3 or decision-material/high-stakes Tier 2. Non-default weighting MUST NOT weaken the base-profile non-dilution rule or hide a material catastrophe cell.

Reference semantics: apply the Section 8.4 TRC loss-reference rule. Scenario losses are residual safe-corridor-relative Base-stream losses for the declared exposure domain; a diagnostic incremental comparison against continuation is not a replacement gating loss.

D.7 Discrete CVaR Algorithm (Canonical; Tier 1-3)

Inputs: losses $L(a, s_1), \dots, L(a, s_n)$; probabilities $p_1, \dots, p_n$; tail level α .

1. Let $b_tail := 1 - \alpha$.
2. Sort scenarios by loss descending: $L_{(1)} \geq L_{(2)} \geq \dots \geq L_{(n)}$, carrying the aligned probabilities $p_{(i)}$.
3. Let k^* be the smallest index such that $\sum_{i=1}^{k^*} p_{(i)} \geq b_tail$.
4. Define $P_{\{k^*-1\}} := \sum_{i=1}^{k^*-1} p_{(i)}$ (and $P_0 := 0$).
5. Define $m_boundary := b_tail - P_{\{k^*-1\}}$.
6. Compute:

$$CVaR_\alpha[L(a)] = (1/b_tail) \cdot \left(\sum_{i=1}^{k^*-1} p_{(i)} L_{(i)} + m_boundary \cdot L_{(k^*)} \right).$$

Audit: PCC includes the sorted table, k^* , $P_{\{k^*-1\}}$, $m_boundary$, and $CVaR_\alpha$.

(Here $L_{(i)}$ is shorthand for $L(a, s_{(i)})$.)

D.8 TRC Corridor Check (Canonical)

$TRC(a) = \text{TRUE}$ if and only if $CVaR_\alpha[L(a)] \leq \tau_TRC$, every active class-specific cap or tolerance under Section 8.4A is satisfied, no unresolved material C3 or C4 field remains, and all mandatory scenario-family requirements are satisfied or governed exceptions are recorded.

D.9 Default TRC Parameters by Context (Canonical Defaults)

Context	α (tail level)	τ_TRC (threshold)
Personal	0.90	0.30
Organizational	0.95	0.20
Reversible policy	0.95	0.15
Irreversible policy	0.99	0.10
Existential risk	0.999	0.05

D.10 Tier 2 Strict TRC Rule (Canonical)

If catastrophe relevance is plausible (per Section 8.1.1), Tier 2 MUST:

- Compute TRC using this appendix
- Use ≥ 5 scenarios minimum
- Include mandatory tail categories with $p_{\text{floor}} \geq 0.02$ per category unless explicitly implausible with documented justification per Section 8.3.2
- Record the full scenario table and CVaR computation in PCC

D.11 Tail Emergency Mode Reference Procedure (Canonical)

If A_{NCRC} is non-empty but no option passes TRC, the ordinary admissible/selectable set is empty. Implementations MUST NOT rank the failed options through ordinary RLS and MUST NOT automatically select the least-CVaR option.

Default disposition: $\text{TRC_ALL_FAIL} \rightarrow$ redesign, delay, no-action comparison, escalation, or refusal.

A provisional action may be considered only through the Tail Emergency Mode requirements in Section 8.8. The PCC MUST demonstrate necessity, failed or unavailable safer alternatives, explicit no-action/delay comparison, independent challenge, an absolute maximum exposure bound, minimized scope and duration, monitoring and shutoff, mitigation, remedy/redress, higher approval, expiry, and return-to-normal criteria.

Only candidates satisfying every Section 8.8 prerequisite may enter the emergency candidate set. Within that set, minimum CVaR may identify the least-tail-risk provisional action. Record

$\text{TAIL_EMERGENCY_MODE_INVOKED}$ and return $\text{TAIL_EMERGENCY_PROVISIONAL_ACTION}$.

This state is not TRC_PASS , ordinary admissibility, ordinary selectability, alignment, or proof of safety. Ordinary RLS is prohibited. If no candidate satisfies the absolute exposure boundary or the governance prerequisites, the result is $\text{TAIL_EMERGENCY_REFUSED}$ or escalation under the governing charter or public-law emergency authority.

End Appendix D.

APPENDIX E: UCI OPERATIONALIZATION PACK (TIER 1-3)

Appendix E is normative for the UCI interface and Tier 3 constraints; indicator families are canonical guidance unless a deployment supplies validated instruments.

Operational maturity note (Normative clarification). Tier-3 conformance requires declared structural indicator families and structural independence; it does not imply universal calibration. Deployments SHOULD publish a provisional indicator companion pack when using local instruments, including sources, normalization rule, and basic reliability notes.

E.1 UCI Component Definitions (Canonical)

Component	Symbol	Description
Cohesion	H_u	Internal connectivity, trust, shared identity, conflict resolution capacity
Flow	F_u	Coordination throughput, information fidelity, resource allocation efficiency
Resilience	R_u	Redundancy, robustness, recovery speed, adaptive capacity
Equity	E_u	Fair distribution of burdens/benefits, voice representation, inclusion

E.2 UCI Formula (Canonical)

Let A_u be the set of applicable UCI components for Union Scope u . The canonical formula is:

$$UCI_u = [\sum_{f_j \in A_u} \alpha_j X_{\{u,j\}}] / [\sum_{f_j \in A_u} \alpha_j]$$

Default: $\alpha_H = \alpha_F = \alpha_R = \alpha_E = 0.25$ before applicable-component renormalization. NA components are excluded; UNKNOWN components are not scored and make UCI unavailable for the affected claim. For Self (U1), Equity is ordinarily NA unless a governed within-self equity instrument is declared. For Biosphere (U7), Equity is ordinarily NA unless a governed ecological-distribution or ecological-justice instrument exists. No perfect, neutral, or identity value is inserted.

E.3 Structural Independence Rule (Tier 3 Binding)

UCI MUST be computed from structural/process indicators distinct from welfare indicators used for RLS.

At Tier 3:

- Deriving UCI from welfare-cell impacts is PROHIBITED
- If structural indicators are unavailable, UCI is treated as unavailable (see E.7)

E.4 Indicator Families by Union (Canonical Guidance)

U_1 Self:

- Cohesion: Psychological integration, goal coherence, self-trust, internal conflict resolution
- Flow: Task execution reliability, attention stability, cognitive throughput proxies
- Resilience: Stress recovery, adaptive coping, flexibility under change
- Equity: ordinarily NA unless a governed within-self equity instrument is declared; if applicable but unmeasured, UNKNOWN

U_2 Household:

- Cohesion: Relationship quality, conflict resolution, trust among members
- Flow: Resource pooling efficiency, coordination routines, decision-making speed
- Resilience: Emergency preparedness, support redundancy, recovery from shocks
- Equity: Fair burden-sharing, voice parity, caregiving distribution fairness

U_3 Community:

- Cohesion: Social capital, trust indices, network density, belonging measures
- Flow: Collective action capacity, coordination lag, information spread fidelity
- Resilience: Mutual aid redundancy, disaster response capacity, recovery history
- Equity: Inclusion of marginalized groups, access parity, procedural fairness in local governance

U_4 Organization:

- Cohesion: Culture trust indices, turnover stability, safety culture, shared purpose
- Flow: Process throughput, coordination efficiency, error rates, execution reliability
- Resilience: Redundancy, continuity planning, incident response maturity
- Equity: Pay fairness, promotion parity, grievance procedures, representation

U_5 Polity:

- Cohesion: Institutional trust, social cohesion indices, legitimacy measures
- Flow: Governance effectiveness, service delivery reliability, coordination speed
- Resilience: Crisis response capacity, redundancy of critical systems
- Equity: Rule of law parity, civil rights access parity, representation integrity

U_6 Humanity / Global Coordination:

- Cohesion: Cross-polity cooperation capacity, treaty adherence norms

- Flow: Global coordination throughput, information-sharing integrity
- Resilience: Global response capacity to pandemics/climate/conflict
- Equity: Burden-sharing fairness, inclusion of vulnerable polities/populations

U₇ Biosphere:

- Cohesion: Ecosystem connectivity, biodiversity integrity, trophic network stability
- Flow: Nutrient cycling integrity, carbon sequestration capacity, hydrological cycle stability
- Resilience: Recovery capacity, redundancy of functional species, anti-fragility markers
- Equity: Use distributional ecosystem-integrity or ecological-justice indicators when a governed instrument exists. If U7 Equity is genuinely inapplicable to the declared structural question, record E7 = NA, exclude it from the UCI aggregation, and renormalize the remaining applicable components. If U7 Equity is applicable but not adequately measured, record E7 = UNKNOWN; do not insert 1.0 or another neutral value. Material UNKNOWN status makes UCI unavailable for the affected claim and requires evidence collection, claim narrowing, redesign, escalation, or refusal of an unqualified CSV pass.

E.5 ΔUCI Computation (Canonical)

$$\Delta UCI_u(a) = UCI_u(a) - UCI_u(\text{baseline})$$

E.6 Prospective UCI Estimation (Tier 3) (Normative Constraints)

Tier 3 requires prospective (ex ante) estimation of UCI changes from structural indicators, not from welfare impacts. The PCC must:

- Identify structural indicators used
- Specify baseline values
- Specify predicted changes under each option
- Specify normalization to [0,1] for each component
- Compute UCI_u and ΔUCI_u using E.2 and E.5

E.7 UCI Unavailability Rule (Tier 3) (Normative)

If structural indicators are unavailable such that UCI cannot be computed without violating E.3:

- UCI MUST be treated as unavailable for tie-break purposes
- If RLS lead is non-decisive, decision MUST escalate for more data/higher tier OR record a judgment call with label JUDGMENT_CALL_UCI_UNAVAILABLE and a monitoring plan
- Any welfare-derived UCI proxy MUST NOT be used to claim Tier 3 compliance

E.7A UCI/HOI Dual-Use Reporting Interface (Normative for Tier 3)

Purpose. UCI and HOI may serve as structural evidence in two distinct authority roles. Tier 3 implementations MUST record which role is being used.

Required fields: structural_indicator_id; indicator_family; containing_scope_affected; evidence_status; measurement_readiness_level; containment_relevance; containment_threshold_declared; containment_disposition; residual_tiebreak_relevance; duplicate_use_note; reviewer_disposition; audit_flag_if_any.

Rule. If containment_relevance is REQUIRED or THRESHOLD_CROSSED, the indicator SHALL be evaluated inside CSV before RLS. If the same indicator is later cited in residual tie-break, the PCC MUST state that gate-level containment disposition was already resolved and that residual use is limited to non-decisive structural preference, monitoring, or hollowing documentation.

Prohibited use. An implementation SHALL NOT move a gate-relevant structural indicator to residual tie-break in order to avoid containment failure, evidence unavailability, or escalation.

E.8 Minimal Measurement Protocol (Canonical Guidance)

For each UCI component, the PCC SHOULD document: indicator(s) used; data source; normalization method; expected directionality; uncertainty or reliability notes; and review cadence.

- Indicator(s) used
- Data source
- Normalization method
- Expected directionality
- Uncertainty or reliability notes
- Review cadence

E.8A Provisional UCI Measurement Protocol for Tier 3 pilots (Canonical Guidance)

When a deployment claims Tier 3 run-level conformance and uses UCI operationally, each UCI component SHOULD be built from one to three structurally independent indicators per relevant scope, normalized to [0,1] under a declared rule, and aggregated by a declared within-component operator (default arithmetic mean unless another operator is justified).

Staleness and failure handling. If a required structural indicator is stale, contested, or unavailable such that structural independence cannot be preserved, the affected UCI component SHALL be treated as unavailable and the run SHALL follow the unavailability rule in E.7 rather than substituting welfare-derived proxies.

Status note (Normative clarification). Tier 3 run-level conformance claims are distinct from framework-level empirical operational-readiness claims. The current release line specifies how provisional UCI constructions are to be documented and constrained, but it does not claim that universally validated UCI instruments already exist.

E.8B Tier 3 Starter Indicator Pack (Canonical Guidance)

The following starter indicators MAY be used as provisional structural indicator families when no validated domain-specific pack exists. They do not constitute final empirical calibration, but they provide a minimum auditable starter surface.

U3 Community

Cohesion: local trust survey or equivalent community-trust instrument

Flow: coordination-response latency or meeting-to-action completion rate

Resilience: mutual-aid activation readiness or recovery-from-disruption measure

Equity: access-to-participation parity or grievance-resolution parity

U4 Organization

Cohesion: psychological-safety / trust instrument or equivalent (e.g., Edmondson, 1999)

Flow: process throughput or median handoff completion time

Resilience: incident recovery time or continuity readiness metric

Equity: promotion / burden / grievance parity indicator

U5 Polity

Cohesion: institutional trust or legitimacy measure

Flow: service-delivery or review-cycle timeliness measure

Resilience: emergency response readiness or continuity-of-governance measure

Equity: appeal access parity or due-process access parity measure

Minimum handling rules

Normalization to [0,1] MUST be declared in the PCC.

Data older than 12 months MUST be marked stale unless a stricter domain rule is declared.

If more than one required indicator family for a used scope is missing or stale, the affected UCI component SHALL be treated as unavailable.

A worked example using this starter pack SHOULD be included in any Tier 3 deployment companion pack.

E.9 Measurement Guidance Expansion (Canonical Guidance)

For publication and Tier 3 preparation, each deployment SHOULD maintain a short UCI measurement note per component and scope covering: operational definition, indicator family chosen, normalization rule, expected update cadence, minimum acceptable data quality, independence check, and failure handling when indicators are missing or stale. Where validated instruments do not yet exist, the PCC MUST label the construction as provisional and MUST NOT present the resulting UCI value as empirically calibrated. This guidance is intended to standardize measurement practice without pretending that UCI instruments are already universally validated.

End Appendix E.

APPENDIX F: FAILURE MODES AND ANTI-GAMING CONTROLS (v12.0)

Appendix F is normative for identified gaming vectors and required mitigations. Informative note. This appendix is carried forward substantively unchanged from the v9.0 line unless otherwise stated.

F.1 Identified Gaming Vectors (Canonical)

Vector	Description	Structural Mitigation
Option set manipulation	Excluding viable alternatives so preferred option "wins"	Minimum option requirement (≥ 2); independent challenger in emergencies; document option generation process
Masking abuse	Masking unfavorable cells to inflate RLS	Non-maskable cells (rights + catastrophe) must remain unmasked; audit flags invalidate PCC
Subgroup erasure	Averaging away harms to vulnerable groups	Worst-off subgroup operator for rights; subgroup enumeration; conservative bound when infeasible
Scenario omission	Leaving out unfavorable tail scenarios	Mandatory tail categories and probability floors; scenario count minima; implausibility test per Section 8.3.2
Probability gaming	Assigning implausibly small p_s to tails	Probability provenance required; p_{floor} per category; sensitivity perturbations at Tier 3
Kernel capture	Tweaking kernel edges to favor preferred option	KQS bands; sensitivity perturbations; evidence classes; humility fallback to NONE
Weight capture	Steering HDW to underweight certain unions/dimensions	Constitutional floors; supermajority locks near floors; transparency ledger; conflict disclosure
Emergency abuse	Declaring emergency to bypass rights	Challenger requirement; remediation plan; public disclosure/audit triggers
Confidence inflation	Overstating confidence to increase impact weight	c_k bounded [0.1, 1]; require evidence notes; sensitivity recommended

Vector	Description	Structural Mitigation
Horizon manipulation	Choosing time horizons strategically	Logarithmic temporal weighting $\tau(t)$; declare T_{ref} ; challenge horizon assumptions
Unknown-as-neutral	Leaving cells blank to avoid negatives	Ignorance penalty phantom instance rule; audit flag if non-maskable cell missing

F.2 Structural Anti-Gaming Features (Canonical)

Non-Compensatory Architecture:

- NCRC cannot be overridden by welfare gains
- TRC cannot be overridden by expected value
- Containment cannot be bypassed by RLS optimization

Transparency Requirements:

- All parameters recorded in PCC
- Registry hashing prevents silent modification (Tier 3+)
- Public rationale (5SPR) required

Audit Infrastructure:

- Random audit lottery (recommended 5% of PCCs at Tier 3)
- Red-team protocols for high-stakes decisions
- Configuration drift detection

Learning Accountability:

- NCAR loop tracks prediction accuracy
- Systematic errors trigger parameter review
- Historical PCCs preserved for accountability

F.3 Audit Flag Taxonomy (Informative pointer; non-authoritative)

The human-readable audit requirements are in Section 14.3A, including its typed completion table; Appendix H.3 specifies registry ownership and binding. H.3 intentionally contains no hand-maintained duplicate table. Use the exact manifest-pinned machine registry when available for a registry-conformant implementation; this informative pointer does not retire a controlling token or establish external registry parity.

F.4 Red-Team and Audit Lottery (Tier 3 Recommended)

Deployments SHOULD adopt:

- Random audit lottery: for example, 5% of PCCs selected for independent review
- Red-team scenario additions: Adversarial scenario injection for TRC testing
- Kernel perturbation stress tests: Systematic edge perturbation

F.5 Boundary and Threshold Gaming (Normative anti-gaming clarification)

RippleLogic does not assume adversaries will only game welfare scores. Because NCRC, TRC, CSV, and PLSS create hard feasibility and escalation conditions, adversarial pressure may shift toward subgroup definitions, rights-threshold placement, scenario-library construction, probability assignment, containment-scope framing, PLSS local-scope declarations, evidence-sufficiency claims, and maturity classifications.

PCC reviewers SHALL inspect these boundary surfaces for gerrymandering, selective omission, motivated uncertainty, threshold laundering, scope narrowing, and evidence laundering. Where any of these surfaces is material to a gate verdict, escalation, or refusal decision, the run SHALL disclose the contested boundary choice, the evidentiary basis, reviewer disagreement if any, and the sensitivity of the verdict to plausible alternative boundary definitions.

Goodhart pressure is not eliminated by RippleLogic; it is relocated and exposed. The system becomes stronger only when boundary definitions, threshold settings, probability libraries, stakeholder maps, and evidence sufficiency judgments are inspectable, contestable, and preserved in the PCC or companion audit record.

End Appendix F.

APPENDIX G: SGP v8.8 INTEGRATION BINDING (NORMATIVE INTERFACE)

Appendix G is normative for the RippleLogic-to-SGP interface.

G.0 Interim Precautionary Trigger

If credible SGP, taxon/class, stakeholder, or artificial-system evidence identifies a plausible moral patient before formal rights-coverage expansion is complete, the run MUST record

PRECAUTIONARY_PROTECTION_REVIEW. Absence of an existing C_r mapping MUST NOT be treated as permission for harmful treatment.

G.1 Version and Artifact Pin

SGP Version Pin: RippleLogic Canon v13.0 pins Sentience Gradient Protocol v8.8.

ProofPack dependency rule: a replayable package MUST bundle or hash-pin the exact governing SGP source and any relied-upon MPS/FPP/GPR/SPR/ICP/RMCP record. This package supplies SGP_v8.8.docx and its local reading mirrors. SGP-aware test claims remain restricted to explicitly supplied interfaces and fixtures; the documents and scoped tests do not establish Tier-4 ProofPack or external repository conformance.

Legacy rule: v7.x SG_norm, SG_patient_norm, SG_measured_norm, A/B/C pillar values, and Plateau/CMIU fields are retained only as version-pinned historical records. They MUST NOT be translated into v8.8 outputs without re-evaluation. v8.0 MPS/FPP/GPR/SPR/ICP records remain structurally compatible only where the target boundary and evidence record remain valid; RMCP and P100 require an explicit v8.8 evaluation.

G.2 Interface-Visible Outputs

A valid SGP v8.8 record may expose:

- MPS evidence-family levels V, I, T, F, and S, including NE fields;
- MPS status MPS-NE or evidence band MPS-0 through MPS-4;
- MPS evidence band/interval, record validity, process confidence, and governed welfare-inclusion hypothesis set;
- protection posture and typed protection bundle;
- FPP(E) in {0,1};
- GPR(E,r,t) for a named process or institution;
- SPR(E,r,t) for a named consequential role;
- optional ICP(E,D,t);
- optional RMCP(C,D,t), including the dimension vector, uncertainty intervals, boundary/scaffolding record, capacity-performance separation, causal scale, and P100 status where assigned;
- target identity/boundary, evidence trace, dependence graph, alternative explanations, reviewer status, validity status, and revision trigger.

RippleLogic consumes these values exactly as supplied by the valid pinned record. It MUST NOT reconstruct MPS from deprecated pillars, infer FPP from a high band, infer sentience or moral worth from RMCP/P100, or infer authority from MPS/FPP/GPR/ICP/RMCP.

G.3 MPS Evidence-Support Map and Welfare-Hypothesis Interface

SGP output	Evidence-support status	Required RippleLogic posture
MPS-NE	NOT_EVALUABLE	Missing, poorly observable, or indeterminate evidence. Never zero. If decision-material, test both $h=0$ and $h=1$ and apply precaution.
MPS-0	[0.00,0.19]	Admissible evidence currently weighs against welfare-bearing experience. Primary exclusion may be tested, but high-harm or model-sensitive cases retain $h=1$ sensitivity.
MPS-1	[0.20,0.39]	Realistic possibility; include $h=1$ and apply precaution.
MPS-2	[0.40,0.59]	Credible possibility; include $h=1$ with robust protection.
MPS-3	[0.60,0.79]	Substantial evidence; $h=1$ is the primary welfare posture.
MPS-4	[0.80,0.99]	Strong convergent evidence; $h=1$ and presumptive full protection/FPP review.
FPP	N/A: separate protection status	Full non-downgradable rights-of-protection; $H_E=\{1\}$, $h=1$.

The numeric intervals are bounded evidence-support interfaces for governance and review. They are not probabilities, quantities of consciousness, intrinsic-worth measures, or direct welfare multipliers. RippleLogic constructs and reports the required welfare-inclusion hypotheses; it does not convert the MPS interval into a coefficient. The FPP row is excluded from evidence intervals: it fixes protection and welfare inclusion even when MPS is NOT_EVALUATED.

G.4 Human FPP and General FPP Recognition

For every human person H , $FPP(H)=1$ and the welfare-inclusion coefficient is $h=1$. RippleLogic SHALL NOT request or use performance measurement to weaken this status.

For a non-human individual or class, FPP exists only when explicitly assigned by a valid SGP v8.8 governed record under its identity or membership rule. MPS-4 alone triggers FPP review; it does not automatically create FPP. FPP governs protection only and does not create CMIU membership, GPR, SPR, legal personhood, or authority.

G.5 Permitted Welfare-Stream Use

For each affected stakeholder instance k :

- Base stream: $h_{\{k,base\}}=1.0$.
- Welfare stream: for each declared sensitivity evaluation j , choose $h_{\{k,welfare\}}^{(j)}$ from the governed welfare-inclusion hypothesis set for E_k ; do not insert the MPS interval directly as a multiplier.
- Human/FPP rule: $h_{\{k,welfare\}}^{(j)}=1.0$ for every human person and FPP entity.
- For MPS-NE, test $h=0$ and $h=1$ whenever decision-material.
- For MPS-1 through MPS-4, include $h=1$; include $h=0$ only where exclusion remains evidentially admissible and could alter the result.
- A representative display value MAY be shown only when every decision-material hypothesis result and selection sensitivity are also reported.

SGP outputs MUST NOT alter RF/NCRC, TRC, or CSV. Protection is enforced through rights coverage and protection posture, not discounted into a welfare multiplier.

G.6 Participation, Reality-Management Capacity, Power, and CMIU Non-Inference

GPR may guide direct, assisted, or represented participation in a named process. It does not determine moral worth or executive power.

RMCP may describe integrated reality-management capacity and P100 status. Humanity collectively is the provisional current reference calibration anchor, but P100 is open across substrates and does not establish wisdom, FPP, membership, or authority. Capacity and actual stewardship performance **MUST** be reported separately.

SPR is role-specific and may inform an AuthoritySelectionRecord. It does not by itself grant authority and cannot bypass lawful mandate or the cascade.

CMIU is defined by Canon Section 12.7 as a governance-coordination view. SGP does not grant CMIU standing. FPP, MPS, RMCP, or P100 alone never creates membership; GPR supplies participation evidence, and consequential roles require SPR plus the Canon's authority controls.

G.7 PCC Logging and Record Validity Floor

When SGP affects protection, stakeholder inclusion, participation, power review, or Welfare-stream analysis, the PCC **MUST** record:

- SGP version and exact artifact hash/reference;
- target identifier and system/class boundary;
- record validity status;
- MPS band or MPS-NE status, evidence-support interval where applicable, confidence, protection posture, and welfare-inclusion hypotheses evaluated;
- FPP status and assignment source;
- the welfare-inclusion hypotheses evaluated and sensitivity result;
- evidence maturity, dependence clusters, access limits, and alternative explanations;
- reviewer panel, conflicts, minority opinion, and welfare advocate where required;
- GPR process/band and participation mode where used;
- SPR role/verdict, mandate, limits, responsibility chain, and revocation method where used;
- RMCP target/collective boundary, dimension vector and uncertainty intervals, external scaffolding, causal scale, capacity-performance separation, P100 decision and challenger review where used;
- re-evaluation date and revision trigger.

Validity floor: a record lacking target boundary, version pin, evidence trace, MPS family results, band/interval, confidence, protection posture, reviewer status, or revision trigger is **INCOMPLETE**. An undefined target, unverifiable evidence trace, unethical method, or unpinned version is **INADMISSIBLE**. Neither may be converted to MPS-0. If a run claims SGP-aware weighting but lacks the valid record and interval actually used, the PCC is **INVALID**.

G.8 Rights-Bundle Concordance

The SGP typed protection bundle maps to the Canon's eight compact protected domains without collapsing them:

- welfare protection -> LIFE, BODY, NEED, DIGN;
- substrate integrity -> BODY, LIFE, DIGN;
- cognitive/memory integrity -> INFO, DIGN, LBTY, BODY;
- continuity/survival -> LIFE, BODY, DIGN;
- anti-ownership/exploitation -> LBTY, DIGN, BODY;
- consent/assent/represented interests -> LBTY, PROC, DIGN;
- representation/appeal -> PROC, INFO, DIGN;
- environmental fit -> ECOL, NEED, LIFE, BODY.

A case may engage several domains. Implementations **MUST** preserve the full applicable set rather than selecting one convenient label.

Single-source rule. Appendix G.8 is the controlling current-package concordance. SGP v8.8 Appendix F.2 is a synchronized reading mirror and **MUST** reference this table rather than redefine a different mapping. Any

future change requires a manifest-listed update to both artifacts and a conformance test. SGP v8.8's non-consumption and non-resource-conversion rule for FPP entities is interpreted through these protected domains and the applicable lawful, emergency, consent, and post-mortem constraints; it is not a licence to infer that every living system is an FPP moral patient.

G.9 Conformance Cases

The authoritative interface cases are SGP v8.8 Appendix I and Canon Appendix R. At minimum, validators must cover: human FPP; humanity collective RMCP-P100 without individual-capacity inflation; high MPS with low GPR; high ICP or RMCP with low MPS; high MPS/FPP with low RMCP; open cross-substrate P100 admission; P100 without CMIU/authority; MPS hypothesis-sensitive ranking; FPP without CMIU/authority; missing interval/record refusal; incomplete registry invalidity; legacy v7 non-equivalence; class-prior protection floor; and temporary dormancy/continuity.

End Appendix G.

APPENDIX H: PCC TEMPLATE AND AUDIT FLAGS CANON PACK (TIER 2-3)

H.0.4 No Silent Remainder Rule (Normative)

Required fields may be unknown, estimated, assumed, not applicable, or blank by design, but they may not be silently blank. A required field in a claim-bearing run MUST carry one of the following dispositions: VALUE, ESTIMATE, ASSUMPTION, DECLARED_UNKNOWN, NOT_APPLICABLE, BLANK_BY_DESIGN, or MISSING_UNDECLARED.

MISSING_UNDECLARED on any required field creates SILENT_REMAINDER_INVALID with severity INVALID. DECLARED_UNKNOWN on a gate-critical field, including rights-floor cells, catastrophe-relevant cells, material Containment fields, SVC subconditions, authority-selection fields, or public conformance fields, requires claim narrowing, evidence collection, escalation, or refusal. It MUST NOT default to pass or zero.

The PCC SHOULD record silent_remainder_count. Claim-bearing Tier 2 and Tier 3 runs SHOULD require silent_remainder_count = 0 before public conformance, deterministic-selection, deployment-readiness, or ProofPack-prep claims are made.

H.0 Carried-forward v10.8 Refusal, Closure, Category-Collapse, and CBI Fields

H.0.1 RefusalRecord field (Normative)

Field name: PCC.RefusalRecord

Required when: the run enters Refusal under Section 4.9.

Minimum fields: refusal_id; trigger_code R-1, R-2, R-3, R-4, R-5, or R-6; trigger_description; options_affected; scopes_affected; dimensions_affected; rights_cells_affected; catastrophe_cells_affected; containment_indicators_affected; evidence_basis; provenance_links; no_action_baseline_status; escalation_route; provisional_action_status; re_run_conditions; reviewer_signature_or_role.

Validation rule. If Refusal is triggered and PCC.RefusalRecord is absent, the PCC is INVALID.

H.0.2 ComputationalClosure field (Normative)

Field name: PCC.ComputationalClosure

Required for: all Tier 2 and Tier 3 runs; recommended for Tier 1.

Minimum fields: `scopes_evaluated`; `scopes_guardrail_only`; `scopes_outside_computation`; `time_horizon_evaluated`; `time_horizon_outside_computation`; `kernel_mode`; `propagation_edges_included`; `propagation_edges_excluded`; `externalities_excluded`; `claim_boundary`.

Validation rule. If the run makes claims beyond the declared `ComputationalClosure` boundary, set `audit_flag CLAIM_EXCEEDS_CLOSURE_BOUNDARY` and mark the PCC `INVALID` for that claim.

H.0.2A RealityGroundingRecord field (Normative)

Field name: `PCC.RealityGroundingRecord`

Required for: Tier 2 runs where any stakeholder-discovery, rights, tail-risk, containment, or authority-selection trigger holds; all Tier 3 runs; and any claim-bearing run that depends on a contested material reality claim.

Minimum fields: `RealitySurfaceDescription`; `MaterialRealityClaims`; `EvidenceTrace`; `DeclaredUnknowns`; `TransitionOrActionBoundary`; `ConsequencePathways`; `ClaimBoundary`; `RealityGroundingStatus`; `FalsificationOrRevisionTrigger`; `MonitoringPath`; `VehicleContentEffectNotes` where social or institutional constructs are material; `SurfaceTerritoryLimitations`; `ReviewerNotes`. The first six substantive outputs **MUST** remain semantically aligned with Section 2.1B: reality surface, evidence trace, material unknowns, transition or action boundary, consequence pathways, and claim boundary.

Validation rule. If a material claim has `INSUFFICIENT_GROUNDING` (or legacy `INSUFFICIENT_CONTACT`) and affects NCRC, TRC, CSV, deterministic selection, or authority-selection disclosure, the run **MUST** downgrade the claim, escalate, or refuse the stronger decision state. If `PCC.RealityGroundingRecord` is missing when required, set `audit_flag REALITY_GROUNDING_RECORD_MISSING` (legacy alias `REALITY_CONTACT_RECORD_MISSING`). If declared `RealityGroundingStatus` exceeds what `EvidenceTrace` substantiates, set `audit_flag REALITY_GROUNDING_STATUS_UNSUPPORTED` (legacy alias `REALITY_CONTACT_STATUS_UNSUPPORTED`).

Claim-boundary rule. At Tier 2-3, a `RealityGroundingRecord` provides traceability and reviewer-checkable evidence discipline; it is not itself verified reality grounding. Tier 4/ProofPack claims require `RealityGroundingStatus` to be independently re-derivable from `EvidenceTrace`, referenced evidence objects, and hash-pinned replay materials.

H.0.2B PCC.CategoryGroundingRecord field (Normative)

Field name: `PCC.CategoryGroundingRecord`

Required when: Section 2.1C / Appendix AH trigger conditions hold for a claim-bearing Tier 3 run, high-stakes Tier 2 run, public conformance claim, public deployment claim, deterministic-selection claim, or optimized metric whose category choice is material.

Minimum fields: `material_category`; `category_role`; `primitive_basis`; `definition`; `category_discriminator`; `dependency_position`; `inclusion_boundary`; `exclusion_boundary`; `common_confusions`; `noncollapse_pairs`; `reference_structure`; `evidence_surface`; `source_authority_status`; `category_maturity_status`; `falsification_or_revision_trigger`; `term_refusal_condition`; `downstream_dependencies`; `semantic_debt_flag`; `reviewer_status`; `required_claim_action`; and `category_failure_class` where applicable.

Validation rule. If Category Grounding is required and `PCC.CategoryGroundingRecord` is missing, set `audit_flag CATEGORY_GROUNDING_RECORD_MISSING` and block deterministic framework selection, conformance, public-deployment, and optimized-metric claims until the claim is narrowed, escalated, marked sensitivity-only / monitor-only, or refused. If `category_maturity_status` is `INSUFFICIENT_GROUNDING`, `LEGACY_UNREVIEWED`, or `CONTESTED` beyond the claim boundary, the validator **MUST** reject unqualified deterministic, conformance, public-deployment, and public assurance claims for the affected decision state.

H.0.3 CategoryCollapseSelfTest field (Normative)

Field name: `PCC.CategoryCollapseSelfTest`

Required for: Tier 2 and Tier 3 runs; recommended for Tier 1.

The field SHALL answer each of the following as PASS, FAIL, or IND:

1. Did the run rely on traceability where admissibility was required?
2. Did the run substitute compliance, approval, or certification for correctness?
3. Did the run treat plausibility, fluency, rhetoric, or institutional authority as evidence?
4. Did the run conflate intelligence, capability, or agency with welfare-bearing status?
5. Did the run conflate moral status, protection, or sentience with governance authority?
6. Did the run allow SGP-derived welfare scalars to alter NCRC, TRC, or CSV outcomes?
7. Did the run treat mathematical consistency as empirical validation?
8. Did the run use meta-union, AIU, cosmic, or universal language as a computational override?
9. Did the run use scientific, engineering, legal-regulatory, normative, institutional, certification, compliance, commercial, or anthropomorphic language as warrant for a conclusion in another domain without a declared and supported cross-domain bridge?

Validation rule. Any FAIL SHALL trigger CATEGORY_COLLAPSE_FLAG. If the failure affects admissibility, selectability, or selection, the run SHALL enter Refusal under Section 4.9.

Appendix H is normative for PCC fields and audit flags at Tier 2-3.

H.1 PCC Schema (Human-Readable)

=====

PROVENANCE AND COMPLIANCE CERTIFICATE (PCC)

RippleLogic Framework v13.0

=====

HEADER (REQUIRED Tier 2-3)

Decision ID: [unique identifier]

Timestamp (UTC): [YYYY-MM-DDThh:mm:ssZ]

Decision Owner(s): [names/roles]

Analyst(s): [names/roles]

Reviewer (Tier 3): [name/role]

Spec Version: RippleLogic v13.0

Implementation Tier: [1 | 2 | 3]

Tier Claim Disposition (Tier 2-3 REQUIRED)

PCC.TierClaimDisposition: [TIER_CLAIMED_AND_MET |
TIER_ATTEMPTED_NOT_ACHIEVED_DOWNGRADED]

If PCC.TierClaimDisposition = TIER_ATTEMPTED_NOT_ACHIEVED_DOWNGRADED, PCC MUST specify:

- PCC.TierAttempted: [2 | 3]

- PCC.TierAchieved: [1 | 2]

Tier Downgrade Flags: [TIER_DOWNGRADED_CONTAINMENT = TRUE/FALSE] (required if any gate is unavailable and a Tier 2 downgrade disposition is used).

Constraint (Normative). If TIER_DOWNGRADED_CONTAINMENT = TRUE, then PCC.TierClaimDisposition MUST be TIER_ATTEMPTED_NOT_ACHIEVED_DOWNGRADED and PCC.TierAchieved MUST be 2.

Propagation Mode: [NONE | QUICK]

Kernel Profile: [NONE | profile name + KQS + evidence note]

If Propagation Mode = QUICK, PCC MUST include:

PCC.Kernel.EntryAbsMax: [float in canonical scale]

PCC.Kernel.RowSumAbsMax: [float in canonical scale] (this is $\|K\|_\infty$)

PCC.Kernel.Bounds: [k_max, p_max used]

PCC.Kernel.StabilityPass: [TRUE/FALSE]

Weight Profile: [HDW | PLSS_LOCAL | Declared interim | Uniform]

PCC.BaselineType_Welfare: [STATUS_QUO | FLOOR_REFERENCE | OTHER_DECLARED]

PCC.BaselineType_Gates: [FLOOR_REFERENCE] (MUST be FLOOR_REFERENCE for all cells in C_r u C_cat)

PCC.BaselineReference_Gates: [indicator definitions + target/floor levels + sources + conversion assumptions]

PCC.BaselineFloorReferenceCellsInWelfare: [list of rights/catastrophe cells that remain in welfare scoring under the non-maskable persistence rule and therefore use floor-reference baseline in welfare scoring]

PCC.BaselineDualExceptionBlock: [OPTIONAL structured list by cell or cell-class: welfare baseline used, gate baseline fixed at FLOOR_REFERENCE, rationale, and reviewer sign-off if required]

If PCC.BaselineType_Welfare = OTHER_DECLARED, PCC MUST specify: PCC.BaselineReference_Welfare.

TRC Context Class: [Personal | Organizational | Reversible |

Irreversible | Existential]

Content Hash: [optional Tier 2; required Tier 3; SHA-256]

STEWARDSHIP BLOCK (CONDITIONAL; REQUIRED when Stw_req=1)

PCC.Stewardship.Stw_req: [TRUE | FALSE]

If Stw_req = FALSE, all fields below MAY be omitted.

PCC.Stewardship.DBD.Scope: [list of out-of-scope authority domains]

PCC.Stewardship.DBD.RefusalRule: [refusal/deferral rule identifier]

PCC.Stewardship.APM.AfterActionAuthorship: [0.0–1.0]

PCC.Stewardship.APM.Rationale: [artifact-anchored justification; MUST reference ≥ 2 of 4 APM anchors per Appendix N.4]

PCC.Stewardship.APM.BaselineAuthorship: [0.0–1.0] (OPTIONAL)

PCC.Stewardship.IRB.InterventionCap: [declared caps]

PCC.Stewardship.IRB.Thresholds: [intervene-only-past-harm-point rule]

PCC.Stewardship.InfluenceLedgerRef: [pointer | "NONE" + justification]

PCC.Stewardship.DelayFlag: [TRUE | FALSE] (OPTIONAL)

SCOPE (REQUIRED Tier 2-3)

Decision Question: [clear statement]

Time Horizon: [years]

Primary Affected Scopes (Union Scopes): [list]

Decision Boundary Notes: [what is in/out of scope; why]

Stakeholder Coverage and Mapping (Tier-governed; v13.0)

Stakeholder Coverage Index (SCI): [SCI-0 / SCI-1 / SCI-2 / SCI-3 / SCI-4]

PCC.SCI.BelowMinimumDisposition: [DOWNGRADE_TIER | RERUN_TO_MEET_MINIMUM] (REQUIRED if SCI_BELOW_MINIMUM triggers; see Section 5.0.1 and audit flag SCI_BELOW_MINIMUM).

PCC.SCI.Justification: [text] (REQUIRED if SCI < tier minimum, OR if SCI equals the tier minimum and a Next-Run Upgrade Plan is required).

Scope Coverage Declaration: Mode [FULL_SCOPE / REDUCED_SCOPE]; Active scopes [list]; Omitted scopes [list]; Omission rationales [text]; Blind-spot statement [text]; Escalation triggers that force scope expansion [list].

SCI Next-Run Upgrade Plan (REQUIRED if SCI equals the tier minimum for the claimed tier): [text; what scope expansions or stakeholder mapping improvements will be executed next run; by when].

Admissibility Attestation (REQUIRED if Mode = REDUCED_SCOPE):

PCC.AdmissibilityAttestation.NCRC_full_Cr_computed: [TRUE | FALSE]

PCC.AdmissibilityAttestation.TRC_full_Ccat_computed: [TRUE | FALSE | NA]

PCC.AdmissibilityAttestation.Containment_computed_as_applicable: [TRUE | FALSE | NA]

Notes (Normative): Reduced-scope mode SHALL NOT remove any non-maskable cells or required scope rows from admissibility computation. If any required attestation is FALSE or missing, the run MUST NOT claim tier compliance for the affected admissibility gate(s) and MUST record the relevant audit flag(s).

Stakeholder Instance Map (InstanceMap): Provide a structured list of stakeholder instances with fields: instance_id, description, role/context, exposure pathway, mapped_scope(s), and rationale. Tier 3 REQUIRED, Tier 2 REQUIRED when any SDP trigger holds (Section 2.2A).

Redundancy Handling Declaration (if any multi-scope mapping): Method [EMERGENT_SCALE_JUSTIFICATION / DEDUPLICATION / ALLOCATION]; Effect-token list (if ALLOCATION); allocation coefficients ($\sum_u \alpha_u \leq 1$ per token); disclosure of any deduplication primary scope choices. For EMERGENT_SCALE_JUSTIFICATION, record the distinct scale-specific EffectTokenIDs, evidence links and challenge disposition.

PLSS BLOCK (CONDITIONAL; REQUIRED when PLSS is used)

PCC.PLSS.LSSDeclaration: [final Local Stakeholder Set used for the run, with reference to the instance-to-scope mapping basis]

PCC.PLSS.ProminenceVector: [declared q_u for all scopes, including sub-component basis for magnitude, probability, persistence, and propagation]

PCC.PLSS.Operator: [DEFAULT_GEOMEAN_V1 | OTHER_DECLARED]; if OTHER_DECLARED, record operator, rationale, and sensitivity note

PCC.PLSS.WeightVector: [computed w_u^{PLSS} for all scopes]

PCC.PLSS.ScopePostureClassification: [Prominent | Guardrail for each scope]

PCC.PLSS.EscalationTriggerAssessment: [ET-1 through ET-7 outcome, including any justification for remaining in PLSS posture]

PCC.PLSS.FallbackDeclaration: [uniform residual fallback used if all $q_u = 0$, or other declared fallback with justification]

MEASUREMENT-MATURITY AND DEPLOYMENT-CONTEXT BLOCK (REQUIRED when applicable; Tier 2-3)

PCC.UCI_Measurement_Readiness: [UCI-M0_UNAVAILABLE | UCI-M1_PROVISIONAL | UCI-M2_DOMAIN_SUPPORTED | UCI-M3_VALIDATED]; PCC.UCI_Indicator_Source: [named sources]; PCC.UCI_Normalization_Rule: [declared rule]; PCC.UCI_Structural_Independence_Posture: [independent | correlated-disclosed | unavailable]; PCC.UCI_Reviewer_Reliability_Status: [not-tested | internally-reviewed | IRR-tested | validated].

PCC.Kernel_Use_Posture: [K0_NONE | K1_STARTER | K2_DOMAIN_CALIBRATED | K3_BACKTESTED | K4_REGISTERED]; PCC.Kernel_Validation_Evidence: [none | source note | calibration report | backtest report | registry ID]; PCC.Kernel_Backtest_Status: [not-backtested | partial | complete]; PCC.Kernel_Error_Bounds: [declared bounds or NA].

PCC.HDW_Capture_Audit: [weight source, stakeholder map, participation quality, expert role provenance, conflict disclosures, challenger window, independent review status, captured-context limitation] (required for Tier 3 public or institutional weighting processes).

PCC.PLSS_Escalation_Register: [profile used, local stakeholder set, guardrail scopes, escalation-trigger checklist, non-escalation rationale, externality search method, prominence-sensitivity result, external challenger status, repeated-PLSS-use flag] (required when PLSS is used).

PCC.Deployment_Context_Profile: [LAB | PILOT | INSTITUTIONAL_ADVISORY | INSTITUTIONAL_OPERATIONAL | PUBLIC_POLICY | AGENTIC_EXECUTION | EMERGENCY_TRIAGE] plus evidence-sufficiency, reviewer-independence, audit-visibility, and escalation-threshold notes.

PCC.Biological_Evaluation_Maturity: [BEM-0_NOT_EVALUATED | BEM-1_TAXON_BASELINE | BEM-2_ETHOLOGICAL_GUIDANCE | BEM-3_EXPERT_REVIEWED_BATTERY | BEM-4_VALIDATED_ANNEX] when biological SGP evidence is invoked. Biological measurement immaturity SHALL NOT reduce precautionary protection where credible taxon-level evidence supports moral patienthood.

Option Set O:

Option A: [description]

Option B: [description]

Option C: [description] (if applicable)

Applicability Mask $m(u,d)$ (RLS only):

Masked cells: [list + rationale per cell]

Non-maskable verification:

Rights cells unmasked: [PASS/FAIL]

Catastrophe cells unmasked: [PASS/FAIL]

INPUTS (REQUIRED Tier 2-3)

Uncertainty Method (Tier 2+ REQUIRED for active cells):

PCC.Uncertainty.Method: [A | B | C]

If Method = B:

PCC.Uncertainty.CellConfidenceAggregationMethod: [CCAM_MIN_V1 | CCAM_WMEAN_V1]

PCC.Uncertainty.c_cell_table: [table of $c(u,d,a)$ for active cells, or a retrievable representation sufficient for recomputation]

Weights:

Union weights w_u : [vector or table]

Dimension weights v_d : [vector or table]

HDW parameters (if used): [λ_U , λ_D ; sources]

Rights Canon:

Thresholds θ_r : [table or “current canonical”]

Coverage sets C_r : [table or “current canonical”]

Subgroup policy:

$G_{\{u,d\}}$ method: [enumerated | conservative bound]

γ_{subgroup} (if used): [value]

Notes: [limitations if any]

TRC Inputs (REQUIRED Tier 2 when plausible; REQUIRED Tier 3):

Catastrophe cells C_{cat} : [list]

Catastrophe weights ω : [list]

α : [value]

τ_{TRC} : [value]

Scenario set S :

For each scenario s :

ID, category, narrative, p_s , parameterization notes

Mandatory tails check:

Categories present: [YES/NO per category]

Probability floors met: [YES/NO]

Any implausibility justifications: [YES/NO; if YES, details]

Containment Inputs (REQUIRED Tier 3):

τ_c : [value]

θ_{pos} : [value]

D_c : [value]

UCI method: [indicators used + sources]

If UCI is unavailable for required containing scopes, record `CONTAINMENT_UCI_UNAVAILABLE` and set `PCC.Containment.UCIUnavailableDisposition` to `DOWNGRADE_TIER` or `COLLECT_DATA_RERUN` (and provide a remediation timeline).

NORMATIVE INPUT REGISTER (REQUIRED when material)

`PCC.NormativeInputRegister`: [array; REQUIRED when any welfare-relevant state, rights-relevant state, or enabling condition materially affects admissibility, precautionary handling, selection, or tie-break outcome]

For each entry n in `PCC.NormativeInputRegister`, record:

`NormInputID`: [unique identifier]

`InputType`: [`WELFARE_RELEVANT_STATE` | `RIGHTS_RELEVANT_STATE` | `ENABLING_CONDITION`]

`Materiality`: [`ADMISSIBILITY_MATERIAL` | `PRECAUTION_MATERIAL` | `SELECTION_MATERIAL` | `TIEBREAK_MATERIAL`]

`AffectedObjects`: [cell(s), right(s), scenario(s), gate(s), structural metric(s), or decision element(s)]

`StakeholderOrScopeTarget`: [stakeholder / class / subgroup / instance / scope / scope-set]

EvidenceBasisType: [DIRECT_INDICATOR | IMPACT_INSTANCE_TRACE | RIGHTS_CANON_INTERPRETATION | SGP_INTERFACE | GOVERNED_PROXY | PRECAUTIONARY_DECLARATION | OTHER_DECLARED]

EvidenceReference: [citation / data source / registry ID / PCC pointer / section pointer]

GovernedInterpretationBasis: [NONE | canon section | charter rule | jurisdictional source | instrument identifier | reviewer memorandum]

PrecautionaryBasis: [NONE | explicit uncertainty rationale]

OptionSymmetryStatus: [APPLIES_TO_ALL_OPTIONS | OPTION_SPECIFIC_WITH_JUSTIFICATION]

AnalystReasonCode: [text or controlled token]

SensitivityNote: [REQUIRED if alternative treatment plausibly changes admissibility or selection]

ReviewerNote: [Tier 3 REQUIRED when Materiality = ADMISSIBILITY_MATERIAL, PRECAUTION_MATERIAL, or SELECTION_MATERIAL]

Normative note. The register exists to make reason-for-action inputs auditable when they are not already fully captured by direct impact-instance tables, the canonical rights layer, the canonical catastrophe layer, or declared structural indicators.

IMPACT ESTIMATION (REQUIRED Tier 2-3)

For each option a:

Direct impacts $I_{dir}(u,d,a)$: [summary table or key cells]

Propagated impacts $I_{prop_base}(u,d,a)$ and $I_{prop_welfare}(u,d,a)$: [summary table or key cells]

Key impact instances (at least top contributors):

For each listed instance k:

(u,d) , μ_k , r_k , t_k , ℓ_k , c_k , e_k , $h_{\{k,base\}}$, $h_{\{k,welfare\}}$, rationale/source

e_k disclosure (Normative). If any listed instance uses $e_k \neq 1.00$, the PCC MUST also record the reason code, default-counterfactual, and whether resetting e_k to 1.00 changes admissibility or selection.

ReachBasis (REQUIRED Tier 2+): For each reach term r_k , record the basis and estimator used (allowed bases: population fraction, asset/flow fraction, or area/volume fraction; other bases MUST be explicitly defined), plus source and any conservative bounds.

Missing data:

Any UNKNOWN_IMPACT cells: [list]

Phantom instances applied: [YES/NO]

CASCADE TRACE (REQUIRED Tier 2-3)

NCRC:

For each option: PASS/FAIL

If FAIL: violated rights + $v_r(a)$ values

Emergency Mode invoked: [YES/NO]

Challenger conducted: [YES/NO]

Challenger notes: [summary]

TRC:

Triggered (catastrophe relevance plausible): [YES/NO]

For each option:

CVaR_α: [value]

Corridor τ_{TRC} : [value]

PASS/FAIL

Tail Emergency Mode invoked: [YES/NO]

Containment (Tier 3 required):

For each option in A_{adm}:

U_{pos}(a): [list of unions]

For each u in U_{pos}: min ΔUCI among ancestors: [value]

PASS/FAIL

PCC.Containment.UCIUnavailableDisposition: [DOWNGRADE_TIER | COLLECT_DATA_RERUN | NA]
(REQUIRED if CONTAINMENT_UCI_UNAVAILABLE flag is recorded).

PCC.Containment.UCIUnavailableRemediationTimeline: [text] (REQUIRED if
CONTAINMENT_UCI_UNAVAILABLE flag is recorded).

A_{sel}: [list]

RLS Ranking:

For each option in A_{sel}:

RLS(a): [value]

$\sigma_{RLS}(a)$: [value] (Tier 3 required)

SignedGap(top, each selectable rival): [values by required variant; sole survivor N/A]

δ : [value]

Decisive: [YES/NO]

Tie-break (if non-decisive):

UCI method used: [UCI-A/B/C]

Δ_{UCI} threshold: 0.05

Result: [winner or non-decisive]

HOI flags (if applicable): [notes]

Escalation invoked: [YES/NO]

SELECTION (REQUIRED Tier 2-3)

Selected Option: [option ID]

Selection Rationale: [brief; must reference cascade outcomes]

Monitoring Plan (NCAR):

Metrics: [list]

Review schedule: [date/interval]

Triggers for re-run: [conditions]

SENSITIVITY ANALYSIS (Tier 3 REQUIRED)

Weights perturbation: [summary; selection stable?]

Threshold perturbation: [summary; any near-miss rights?]

Kernel perturbation: [summary; flag fragile?]

Scenario perturbation: [summary; probability sensitivity]

Robustness classification: [Robust | Sensitive | Fragile]

FIVE-SENTENCE PUBLIC RATIONALE (5SPR) (Tier 2-3 REQUIRED)

1. CONTEXT: [What decision was made and why now?]

2. OPTIONS: [What options were considered?]

3. CONSTRAINTS: [What was eliminated by NCRC, TRC, and/or Containment (and why)?]

4. SELECTION: [Why the chosen option won among selectable options?]

5. MONITORING: [What follow-up will be tracked and when?]

AUDIT FLAGS (REQUIRED when triggered)

[List flags + explanations]

SIGNATURES

Analyst: [name, date]

Decision Owner: [name, date]

Reviewer (Tier 3): [name, date]

=====

SOURCE-COUPLING RECORD (CONDITIONAL; REQUIRED when Section 2.1D trigger holds)

PCC.SourceCouplingRecord.claimed_capability: [capability being claimed]

PCC.SourceCouplingRecord.enabling_conditions: [conditions required for the capability to hold]

PCC.SourceCouplingRecord.boundary_conditions: [domain and scope boundaries]

PCC.SourceCouplingRecord.source_evidence: [evidence surface or warrant]

PCC.SourceCouplingRecord.generator_output_distinction: [how generated output is kept distinct from grounded warrant] PCC.SourceCouplingRecord.inherited_assumptions: [assumptions inherited from source, model, data, or substrate] PCC.SourceCouplingRecord.downstream_compensations: [filters, monitoring, governance, or controls used downstream] PCC.SourceCouplingRecord.source_coupling_status:

[SOURCE_COUPLED | SOURCE_PARTIAL | SOURCE_INFERRED | SOURCE_UNKNOWN | SOURCE_CONTESTED | SOURCE_DEBT_RISK | SOURCE_COUPLING_FAILURE] Machine-readable

records MUST use the canonical SourceCouplingStatus enum above. Human-readable aliases such as GROUNDED may be displayed only as labels for SOURCE_COUPLED; they SHALL NOT replace canonical tokens in PCC or validator-facing records. PCC.SourceCouplingRecord.source_debt_flag: [TRUE/FALSE]

PCC.SourceCouplingRecord.falsification_or_recheck_trigger: [trigger]

PCC.SourceCouplingRecord.required_claim_action: [proceed | narrow | control | redesign | escalate | refuse]

PHYSICAL/CAUSAL ADMISSIBILITY EVIDENCE PROFILE (CONDITIONAL; REQUIRED when Section 2.1E trigger holds)

PCC.PhysicalCausalAdmissibilityEvidenceProfile.candidate_generation_source:

[model/planner/controller/person/procedure that generated the candidate]

PCC.PhysicalCausalAdmissibilityEvidenceProfile.physical_or_causal_model_used: [model or basis]

PCC.PhysicalCausalAdmissibilityEvidenceProfile.validity_domain: [declared domain]

PCC.PhysicalCausalAdmissibilityEvidenceProfile.boundary_conditions: [conditions]

PCC.PhysicalCausalAdmissibilityEvidenceProfile.uncertainty_range: [range/class]
 PCC.PhysicalCausalAdmissibilityEvidenceProfile.failure_modes: [failure modes]
 PCC.PhysicalCausalAdmissibilityEvidenceProfile.reversibility_or_irreversibility_boundary: [boundary]
 PCC.PhysicalCausalAdmissibilityEvidenceProfile.verification_simulation_empirical_test_or_expert_warrant: [external warrant]
 PCC.PhysicalCausalAdmissibilityEvidenceProfile.admissibility_warrant_source: [domain method or authority surface that evaluates admissibility]
 PCC.PhysicalCausalAdmissibilityEvidenceProfile.monitoring_and_shutoff_path: [monitoring/shutoff]
 PCC.PhysicalCausalAdmissibilityEvidenceProfile.residual_unknowns: [unknowns]
 PCC.PhysicalCausalAdmissibilityEvidenceProfile.required_claim_action: [proceed | narrow | control | redesign | escalate | refuse]
 PCC.PhysicalCausalAdmissibilityEvidenceProfile.status: [PCAE_SUPPORTED | PCAE_ASSUMPTION_BOUND | PCAE_PARTIAL | PCAE_CONTESTED | PCAE_UNKNOWN | PCAE_VERIFICATION_REQUIRED | PCAE_CONTROL_REQUIRED | PCAE_REDESIGN_REQUIRED | PCAE_REFUSE_OR_BLOCK]
 PCC.PhysicalCausalAdmissibilityEvidenceProfile.physical_execution_claim_status: [PHYSICAL_ADMISSIBILITY_SUPPORTED | GOVERNANCE_PERMISSION_ONLY | PHYSICAL_ADMISSIBILITY_NOT_ESTABLISHED | PHYSICAL_ADMISSIBILITY_CONTRAINDICATED]

METHODOLOGICAL INTEGRITY RECORD (CONDITIONAL; REQUIRED when Section 2.1F or Appendix AN trigger holds)

PCC.MethodologicalIntegrityRecord.claim_or_component: [claim/component]
 PCC.MethodologicalIntegrityRecord.claim_type: [empirical | causal | formal | operational_capability | legal_regulatory | normative | validation | design_target | metaphysical_horizon]
 PCC.MethodologicalIntegrityRecord.claim_domain: [scientific_empirical | engineering_operational | formal_mathematical | legal_regulatory | normative_ethical | institutional_governance | commercial_descriptive | anthropomorphic_social | metaphysical_horizon]
 PCC.MethodologicalIntegrityRecord.warrant_domains: [one or more domains from which support derives]
 PCC.MethodologicalIntegrityRecord.cross_domain_bridges: [one bridge object per material cross-domain relationship: from_warrant_domain; to_claim_domain; bridge_status; bridge_proposition; bridge_evidence_or_rationale; uncertainty; validity_or_jurisdiction_boundary; revision_or_refusal_condition]
 PCC.MethodologicalIntegrityRecord.dependency_position: [where it sits in the method]
 PCC.MethodologicalIntegrityRecord.starting_assumptions: [assumptions]
 PCC.MethodologicalIntegrityRecord.definition_or_operator_used: [definition/operator]
 PCC.MethodologicalIntegrityRecord.necessity_or_alternative_check: [necessity claim and alternatives]
 PCC.MethodologicalIntegrityRecord.evidence_or_test_surface: [evidence/test surface]
 PCC.MethodologicalIntegrityRecord.falsification_or_revision_trigger: [trigger]
 PCC.MethodologicalIntegrityRecord.uncertainty_or_confidence_method: [method]
 PCC.MethodologicalIntegrityRecord.downstream_dependencies: [dependencies]
 PCC.MethodologicalIntegrityRecord.re_derivation_scope_if_changed: [rerun/rederive scope]
 PCC.MethodologicalIntegrityRecord.reviewer_status: [reviewer status]
 PCC.MethodologicalIntegrityRecord.required_claim_action: [proceed | narrow | control | redesign | escalate | refuse]

H.1A PCC.ReferenceStructureRecord (Normative for Tier 3 public conformance claims; recommended for Tier 2 high-stakes runs)

The PCC SHOULD include a ReferenceStructureRecord for Tier 2 high-stakes runs and SHALL include one for Tier 3 public conformance or validation claims. This record does not alter historical NCRC, TRC, CSV, RLS, or UCI/HOI arithmetic. It records the reference structure used to support the claim being made.

Minimum fields:

record_id

decision_id
 option_id_or_set
 claim_scope
 decision_boundary
 reference_object
 active_reference_layers
 invariant_constraints
 causal_pathway_summary
 falsification_conditions
 evidence_sources
 evidence_status_tokens
 measurement_maturity_status
 adaptation_boundary
 authority_boundary
 refusal_triggers
 excluded_references
 reviewer_disposition
 timestamp
 artifact_hashes

Validation rule. If PCC.ReferenceStructureRecord is required and missing, the PCC is incomplete for any claim that the run was reference-structure-audited. Missing ReferenceStructureRecord does not by itself alter cascade arithmetic, but it prevents stronger public conformance, deployment-readiness, or validation claims.

H.2 Minimum Required Fields by Tier (Canonical)

Field	Tier 1	Tier 2	Tier 3
Decision ID	Recommended	REQUIRED	REQUIRED
Timestamp	Recommended	REQUIRED	REQUIRED
Option set	REQUIRED	REQUIRED	REQUIRED
NCRC results	Heuristic ok	REQUIRED	REQUIRED + subgroup policy
TRC results	Required trigger screen; quantitative upon escalation	REQUIRED when catastrophe plausible	REQUIRED
CSV / Containment and Structural Viability	REQUIRED before any ordinary ranking/selection claim: brief qualitative screen, or an evidence-backed CSV_NOT_MATERIAL rationale	REQUIRED proportional result before RLS	REQUIRED full and binding
RLS scores	Optional	REQUIRED	REQUIRED
Uncertainty + discrimination	Optional	REQUIRED (active cells)	REQUIRED (active cells)

Field	Tier 1	Tier 2	Tier 3
Sensitivity analysis	Optional	Recommended	REQUIRED
5SPR	Optional	REQUIRED	REQUIRED
Audit flags	Optional	REQUIRED when triggered	REQUIRED when triggered
Normative Input Register	Optional	REQUIRED when material	REQUIRED when material

H.3 Audit Flag Canon Pack (Canonical)

Section 14.3A and docs/implementation/CANONICAL_AUDIT_FLAG_REGISTRY_v1.1.yaml jointly form the controlling canonical audit-flag registry. Appendix H.3 intentionally contains no hand-maintained duplicate table because omission or drift in a mirror can silently weaken conformance checks.

Implementations, validators, workbooks, and companion standards MUST bind to that joint registry. A token absent from the YAML is local/informative unless the Canon explicitly assigns it controlling status; any explicit controlling assignment creates a release-blocking obligation to add the token to the YAML before a registry-conformant machine release is published. The reading-package boundary in Section 14.3A applies; missing machine material never downgrades a controlling textual requirement.

Stewardship audit tokens remain controlled by Appendix N.7. Agent-runtime event ownership and its mapping to Canon states are controlled by Section 14.3A.1. Neither surface creates a sixth gate or modifies the RG -> RF/NCRC -> TRC -> CSV -> RLS cascade.

H.10 v13.0 PCC Records and Audit Tokens (Normative)

The following records are required when triggered and SHALL be hash-referenced from the PCC. Machine schemas may use equivalent field casing but MUST preserve the meanings and state distinctions.

Record / token	Minimum role
EffectTokenRegistry	Stable effect identity, primary computational home, secondary views, granularity challenge
RightsEffectTokenRecord	Unallocated rights harm and non-dilution disposition
PartitionInvarianceTest	Declared versus challenger representation and reversal result
ProbabilityOrderReview	Likelihood semantics, operator order, double-counting check
AdverseConfidenceCounterfactual	Declared point result and adverse-bound result
TemporalBurdenReview	Trajectory, recurrence, cumulative/terminal/irreversible burden
DistributionalBurdenReview / SubgroupClosureRecord	Worst subgroup, search completeness, CSV routing
ReferenceFamilyRecord / BaselineFamilyRecord	Ex-ante admissible families and ranking results
ComparisonMaskRecord	Run-level mask, active-cell set, justified exceptions

Record / token	Minimum role
GapSensitivityRecord	Independence, dependence/covariance treatment, epsilon effect
WeightProfileRecord	Status, authority basis, competing profiles
OptionClosureRecord	No action, delay, redesign, mitigation, hybrid/challenger disposition
RLSRobustnessRecord	Finite module states and unique-selection eligibility
Key audit tokens	See Section 14.3A and Appendix H.3 for controlling token types, actions and registry binding. This row is a pointer, not an independent enumeration.

H.10.1 Minimum Record Content and Binding (Normative)

A record required by H.10 SHALL provide, directly or through an unambiguous hash-bound PCC reference: stable record identity and type; the owning run, option/comparison and configuration; applicability and trigger rationale; governing source clause(s); evidence/input references; result or explicit unresolved/not-evaluated posture; review/disposition; and links to dependent records. A documented equivalent field casing or container is permitted. Shared metadata need not be copied into every child. Missing information is not zero, an empty reviewed set or PASS. If not triggered, retain the rationale required by the owning rule; if triggered but incomplete, the existing missing-record, unknown and robustness rules control. A hash establishes referenced-byte identity, not evidence truth.

H.10.2 Record-Specific Minimum Content (Normative consolidation)

Existing record family	Minimum content, in addition to H.10.1	Controlling source
EffectTokenRegistry	For each token: stable ID, bearer/class, changed state, option, baseline, contiguous window, scenario and evidence references; one primary home or permitted conserved allocation; secondary views; distinct downstream links; identity/revision and granularity adjudication.	§2.2A.4A; §5.0C; §5.5A
RightsEffectTokenRecord	Bind the rights token to its bearer, protected rights, evidence and any parent effects; retain unallocated severity/exposure and the applicable floor, categorical and severe-hazard channel inputs/results, subgroup denominators, unresolved fields and subsequent allocation links.	§7.3–§7.4.5
PartitionInvarianceTest	Record trigger, declared and evidence-supported challenger representations, invariant facts/evidence/options/weights/parameters, outputs for each representation, any leader or decisiveness reversal, and outcome-independent resolution or refusal.	§5.5A; R.22
ProbabilityOrderReview	Identify affected effects/events, likelihood semantics, scenario-versus-instance probability ownership, and operator order. Record the double-counting check and all selection-material admissible alternative interpretations with their outputs and disposition.	§5.2B
AdverseConfidenceCounterfactual	Identify adverse ordinary-welfare instances and confidences, preserve the declared point construction, specify the $c_k=1$ adverse counterfactual, hold unaffected inputs fixed, and report changed scores, leaders, Gaps, mitigation/control claims and disposition.	§5.2C

Existing record family	Minimum content, in addition to H.10.1	Controlling source
TemporalBurdenReview	Bind effects to their contiguous decision-relevant windows and trajectories; record recurrence, cumulative and terminal burdens, peak severity and irreversible endpoints, required temporal sensitivities, mitigation and routing to rights/TRC/CSV where material.	§5.3–§5.3B
DistributionalBurdenReview	Identify affected subgroups and concentrated burdens, reported profiles and worst-off cases, material dependencies/viability effects, control and capacity evidence, and the resulting existing CSV disposition. It is not a new gate.	§2.4A; §9.8
SubgroupClosureRecord	Record the protected/affected-population boundary, subgroup discovery methods and sources, groups considered, plausible omissions with rationale, material unknowns, review and closure status. An unsearched population is not an empty reviewed set.	§2.4A; §7.3
ReferenceFamilyRecord	Identify the normalized construct and ex-ante reference family, legitimacy/evidence basis, admitted and rejected reference alternatives, fixed comparison inputs, and each alternative's impacts, ranking/decisiveness and sensitivity disposition.	§5.1A; R.27
BaselineFamilyRecord	Identify ex-ante baseline alternatives and their evidence/authority basis; distinguish floor/safe-corridor references from status-quo or continuation comparisons; retain each ranking result and disposition without replacing the controlling gate baseline.	§5.1A; §10.2; R.28
ComparisonMaskRecord	Record the common 49-cell comparison mask, active and non-maskable cells, effective mass Q, and any justified option-specific structural exceptions with evidence and reviewer status; retain the common-basis sensitivity and its result.	§10.2B
GapSensitivityRecord	Record the ranking basis/contenders, final-score uncertainty interpretation and evidence, within-option clusters, cross-option joint-error treatment or justified non-materiality, epsilon/delta and required variants, and every contender's signed comparison, leader and disposition.	§10.3–§10.4B
WeightProfileRecord	Identify w, v, m and kappa profiles, floors, effective weights and Q, legitimate authority/provenance, ex-ante lock, competing legitimate profiles and resulting rankings. Normative disagreement is not converted into statistical uncertainty.	§10.1–§10.2A; §13
OptionClosureRecord	Record the decision boundary, candidate options, no-action/delay/redesign/mitigation/hybrid and challenger paths considered, evidence-supported inclusion/exclusion rationales, independent challenge where required, and remaining option-set limitations.	§15.4B
RLSRobustnessRecord	Use the existing finite module registry and its five states, with triggers/non-trigger reasons, supporting variant records, leader and every-contender results, unresolved modules, and resulting unique-selection eligibility/disposition. Do not substitute a blanket PASS.	§10.4B

H.10.3 Field and View Boundaries (Normative clarification)

HumanCompensationLoadRecord already has its minimum fields in Section 4.10B, and PeakHarmReview has its trigger and fields in Section 5.3. These are not undefined merely because their type names occur once. PCC.CBI.ComputedObject identifies the computed object and interpretation; it is not a new independent record family. ReasonCodeSet is the set of permitted instrument reason codes in PCC.EKRegistry. MaturityProfileRecord refers to the typed claim/evidence-maturity information under Appendix AQ, not one

scalar of readiness. An implementation SHALL bind such references to the owning record and preserve its fields rather than invent a second uncontrolled schema. The minima above specify semantic content; the separately versioned external run-record schema and registries are still required for claims of exact external wire/registry conformance.

APPENDIX I: TIER-4 DESIGN TARGET (NOT CLAIMABLE IN THE CURRENT RELEASE)

I.1 Status Declaration

Tier-4 Design Target (Not Claimable in Any Current Release Line, Including v13.0)

Tier 4 is specified as a target profile requiring artifacts and capabilities that are not yet publicly available. Tier-4 compliance MUST NOT be claimed until ProofPack is publicly released and independently replayable.

Reality-contact Tier-4 requirement. Any future Tier-4 claim that relies on RealityGroundingStatus MUST make that status independently re-derivable from EvidenceTrace and supporting artifacts. Operator-declared status is sufficient for Tier 2 traceability only when review status is explicitly recorded; Tier 3 must follow its declared reviewer regime; neither is sufficient for Tier 4 admissibility, ProofPack readiness, or machine-verifiable ecosystem claims.

I.2 Tier 4 Requirements (Target Specification)

Hash-Bound Registries:

- Rights anchors registry (REG-RIGHTS-ANCHORS-v1)
- Catastrophe indicators registry (AF-BASE)
- Scenario library registry
- Kernel registry with evidence classes
- HDW ballots and weights registry

Deterministic Numeric Profile (NDP_FIXEDPOINT_V1):

- Fixed-point scale: $S = 10^9$
- Representation: $X_{fp} = \text{round_half_even}(x \times S)$ as signed int64
- All divisions use round-half-even
- Hard-fail on overflow with audit_flag NUMERIC_OVERFLOW

Saturation Lookup Table (SAT_LUT_FP_V1):

- Pre-computed tanh values for fixed-point inputs
- Hash-bound in ProofPack
- Runtime computation of tanh PROHIBITED

Temporal Weight Registry (REG_TEMPORAL_WEIGHTS_V1):

- Pre-computed $\tau(t)$ values for standard horizons
- Hash-bound in ProofPack
- Runtime computation of logarithms PROHIBITED

Canonical JSON Profile:

- NO_FLOATS rule: All numeric values as exact rationals {num, den}
- Deterministic key ordering
- Canonical whitespace rules
- SHA-256 hashing for integrity

Propagation Mode Restriction:

- FULL propagation PROHIBITED for Tier 4 Pilot-Executable

- NONE or QUICK only
- FULL available only in future Tier 4 Certified profile

I.3 ProofPack Contents (Target)

The ProofPack artifact bundle will provide:

- Schemas: JSON schemas for PCC and all registries
- Manifests: Hash indices for all artifacts
- Canonicalization Rules: Exact specification for deterministic hashing
- Registries: Machine-readable canonical registries
- Test Vectors: Reference inputs and expected outputs for validation

I.4 Replay Test Requirement

A Tier 4 PCC MUST pass a replay test: independent re-execution using PCC inputs and referenced registries reproduces:

- Identical admissibility outcomes (NCRC, TRC, CSV)
- Identical RLS ranking
- Identical selected option

Failed replay invalidates the Tier 4 claim.

End Appendix I.

APPENDIX J: VALIDATION PROTOCOL PACK

Appendix J is normative for pre-registration templates and scoring definitions when validation claims are made.

J.1 Pre-Registration Template (Normative)

=====

RIPPLELOGIC VALIDATION STUDY PRE-REGISTRATION

=====

STUDY IDENTIFICATION

Study ID: [unique]

Registration date: [YYYY-MM-DD]

Principal investigator: [name, affiliation]

Funding source: [if any]

RippleLogic version: [v13.0]

Tier executed: [1|2|3]

Domain: [AI deployment | policy | org strategy | etc.]

HYPOTHESES

Primary hypothesis: [H1/H2/etc. with exact measurable claim]

Secondary hypotheses: [list]

Falsification criteria: [explicit thresholds from Section 17]

METHODS

Design: [pilot / RCT / quasi-experimental / backtest]
Sample: [population, size, selection method]
Comparator: [baseline governance / alternative framework]
Duration: [time]
Decision types included: [list]
Exclusion criteria: [list]

MEASURES

Primary outcomes: [definition, data source, timing]
Secondary outcomes: [definition, data source, timing]
Rights measurement: [how rights violations are verified]
Tail-loss measurement: [how catastrophic outcomes are defined]

ANALYSIS PLAN

Statistical methods: [tests/models]
Effect size thresholds: [minimum meaningful difference]
Multiple comparison correction: [method]
Missing data handling: [approach]
Sensitivity checks: [weights, scenarios, thresholds]

GOVERNANCE

Ethics approval: [status/ID]
Data management: [storage, access, retention]
Reporting commitment: [timeline and venue]

=====

J.2 Outcome Measures by Union and Dimension (Canonical Examples)

Union	Dimension	Measure Type	Example Indicators
Self	Health	Validated survey / admin	SF-36, PHQ-9, GAD-7, injury rates
Self	Agency	Validated survey	Autonomy/self-efficacy scales
Household	Material	Administrative	Income stability, housing security
Community	Social	Survey / network	Trust indices, social capital measures
Organization	Flow/Resilience (UCI)	Operational	Throughput, incident recovery time
Polity	Agency	Index	Participation rates, rule-of-law indices
Biosphere	Environment	Monitoring	Emissions, biodiversity integrity

J.3 Exact Sign Accuracy and Partial-Credit Diagnostic (Normative)

For each cell in the preregistered evaluation set, define predicted and observed signs using the same declared zero/deadband and observation horizon. Preserve the evaluation unit and case/dependence structure.

Predicted sign: $\text{sign}(\text{l_prop}(u,d,a)) \in \{-1, 0, +1\}$

Observed sign: $\text{sign}(\Delta_{\text{observed}}(u,d)) \in \{-1, 0, +1\}$

ExactMatch = 1 if predicted = observed, and 0 otherwise. ExactSignAccuracy = (sum ExactMatch) / n. In this candidate, SignAcc denotes ExactSignAccuracy; the prior partial-credit definition is renamed PartialSignScore below and cannot be pooled as an unchanged metric across versions.

- For the separate PartialSignScore diagnostic: PartialMatch = 1 if predicted = observed.
- PartialMatch = 0.5 only when signs differ and (predicted = 0 or observed = 0).
- PartialMatch = 0 for all other unequal signs.

PartialSignScore = (sum PartialMatch) / n. It is not exact sign accuracy and has no inherited 0.70 acceptance threshold.

The 0.70 value is retained only as the provisional exact-agreement research hypothesis in Section 17.1; it is not a validated threshold or an automatic acceptance transfer from the legacy partial-credit metric. Confirmatory acceptance requires metric-specific justification and adoption before outcomes are seen. Preregister the no-change and majority baselines, evaluation set, deadband, missing-outcome policy, adverse/beneficial/no-change class measures, changed-cell exact accuracy and uncertainty with case dependence. Report denominators and class counts; empty or undefined measures are unavailable, not perfect. Directional skill requires the preregistered baseline comparison and class evidence, not the target alone. Do not count correlated cells as independent cases or credit pre/post change as an identified treatment effect without Section 17.5 evidence.

J.4 Magnitude Error (RMSE) (Normative)

$\text{RMSE} = \sqrt{[(1/n) \sum_i (\text{l_prop}_{,i} - \text{observed}_{,i})^2]}$

Default target: $\text{RMSE} \leq 0.25$ (domain dependent)

J.5 Backtest Procedure (Normative)

Case selection: Identify historical decisions with measurable outcomes and reconstructable option sets.

1. Information reconstruction: Use only data available at decision time (no hindsight leakage).
1. Blind analysis: The analyst SHOULD be unaware of actual outcomes during the run.
1. Execute RippleLogic run (Tier declared) and emit Backtest PCC.
1. Compare predictions to realized outcomes using declared metrics (for example SignAcc, RMSE, rights detection).
1. Report results and failures; update methods via NCAR with versioning.
1. End Appendix J.

Appendix K is informative and strongly recommended for any deployment using Starter KOPS or provisional QUICK propagation.

APPENDIX K: STARTER KERNEL COMPANION PACK (INFORMATIVE)

K.1 Purpose

Starter KOPS provides an informative, provisional kernel surface for transparency, sensitivity testing, and early implementation support. It MUST NOT be treated as a fully validated propagation model. The Bio-Environment to Humanity-Health and Self-Health to Household-Material starter entries below are SUSPENDED for numerical propagation pending a reviewed sign warrant. Retained coefficients document the prior proposal; they are not usable defaults. A new warrant must state higher-is-better source/target constructs, domain, lag, evidence and causal sign, and test positive and negative source perturbations. Do not silently reverse signs or

treat suspension as evidence of no pathway. Use the existing humility fallback, disclose omitted material pathways, and narrow or refuse claims when they cannot be bounded.

K.2 Evidence Class Summary Table K.1: Evidence Class Summary

Class	Name	Minimum Evidence
A	Validated	Replicated empirical studies with quantified effect sizes and domain match
B	Supported	Multiple studies consistent in sign; some quantification
C	Plausible	Theoretical support with limited empirical evidence
D	Speculative	Expert judgment without empirical support
E	Elicited	Structured elicitation for this framework (lowest confidence)

K.3 Provisional Edge Table Table K.2: Provisional Starter KOPS Edge Table

Source Cell (u,d)	Target Cell (u,d)	k_raw	k_shrunk	Evidence Class	Rationale Family
(4,1) Org-Material	(2,1) HH-Material	0.40	0.14	E	Employment income propagation
(4,1) Org-Material	(3,1) Comm-Material	0.30	0.11	E	Local multiplier effects
(7,7) Bio-Environment	(6,2) Humanity / Global Coordination - Health	-0.35	-0.12	B	SUSPENDED: climate/health sign warrant required
(7,7) Bio-Environment	(6,7) Humanity / Global Coordination - Environment	0.50	0.18	B	Earth system propagation
(5,5) Polity-Agency	(4,5) Org-Agency	0.25	0.09	E	Institutional enabling effects
(3,3) Comm-Social	(1,3) Self-Social	0.45	0.16	B	Social capital / wellbeing
(3,3) Comm-Social	(2,3) HH-Social	0.35	0.12	E	Family-community linkage
(4,2) Org-Health	(1,2) Self-Health	0.40	0.14	B	Occupational health pathway
(6,4) Humanity / Global Coordination - Knowledge	(5,4) Polity-Knowledge	0.30	0.11	E	Diffusion / coordination knowledge
(1,2) Self-Health	(2,1) HH-Material	-0.25	-0.09	B	SUSPENDED: health/earnings sign warrant required

K.4 Mandatory PCC Disclaimer “Ripple propagation used a PROVISIONAL Starter KOPS kernel. Coefficients are not fully validated. Sensitivity analysis was performed per RippleLogic v13.0 requirements; conclusions are contingent on kernel uncertainty.”

End Appendix K.

APPENDIX L: GLOSSARY (v12.0)

Appendix L is informative but standardized terminology is strongly recommended for audit consistency.

Stakeholder: Any entity plausibly affected by a candidate option through direct impact, externalities, indirect pathways, or future propagation within the declared horizon.

Instance: A stakeholder-in-role-in-context bound to a decision, used for auditability and mapping into Union Scopes.

Stakeholder Discovery Protocol (SDP): The tier-linked minimum process for discovering omission-likely stakeholders and recording discovery evidence (Section 2.2A).

Stakeholder Coverage Index (SCI): A tier-governed disclosure ladder describing the stakeholder coverage posture of a run, used for auditability, not for score changes (Section 5.0.1).

InstanceMap: The structured mapping from stakeholder instances to one or more Union Scopes, recorded in the PCC (Appendix H).

Scope Coverage Declaration: A PCC section that declares whether the run is full-scope or reduced-scope, lists omitted scopes and rationales, and specifies escalation triggers (Appendix H).

Active cell: A Scope×Dimension cell that participates in RLS, or is in any rights/catastrophe coverage set, or is populated for any gate or justification, and therefore must not be treated as unknown-by-default (Section 10.3).

ConservedAllocationTokenID (allocation only): A redundant conserved-unit causal token used solely to prevent redundant counting across scopes under the ALLOCATION redundancy method (Section 2.2A).

Reduced-scope mode: A reporting or mapping-depth mode that omits one or more scopes, permitted only with explicit PCC disclosure and without removing any non-maskable cells from admissibility gates (Section 2.2A).

SG_norm / SG_patient_norm / SG_measured_norm: Deprecated v7.x interface fields. They may appear only in pinned legacy records and are not equivalent to the v8.x typed outputs used by current SGP v8.8.

MPS(E,t): The v8.x moral-patienthood evidence surface, current in SGP v8.8, used to determine protection posture and required welfare-inclusion hypotheses. It is not exact moral worth, a consciousness probability, or a cardinal welfare multiplier.

FPP(E): The v8.x full non-downgrade protection status, current in SGP v8.8. For RippleLogic welfare sensitivity, every FPP entity uses the fixed welfare-inclusion set {1}. Any legacy [1,1] display denotes that inclusion set only; it is not an MPS evidence interval.

GPR and SPR: Separate v8.x role surfaces, current in SGP v8.8, for participation design and role-specific consequential-power review. Neither changes MPS or FPP.

Full Protection Plateau (FPP-100): The constitutional rights-of-protection summit defined by SGP. Within RippleLogic it fixes full protection and the welfare-inclusion set {1}; it does not become a prestige score, alter earlier gates, establish CMIU membership, or grant governance authority.

MPS-4 / presumptive-full-protection posture: The strongest evidence band below an assigned FPP status in SGP v8.8. It triggers strong protection and FPP review but is not a scalar percentage of moral worth and does not itself confer FPP.

CMIU: Collective Managing Intelligence Union, a Canon-defined U6 governance-coordination view for managing intelligences and institutions. SGP does not assign CMIU membership. MPS/FPP do not create CMIU membership; GPR informs participation design, and consequential roles require SPR plus lawful governance controls.

Human Full Protection Plateau: The non-overrideable rule that every human person has FPP and the fixed RippleLogic welfare-inclusion coefficient $h=1$, independent of measurement, disability, age, illness, wakefulness, or observability. This is not an MPS evidence interval.

SGP v8.8 Binding Interface: The only permissible interface by which RippleLogic consumes MPS, FPP, GPR, SPR, optional ICP, and informative RMCP/P100. Defined normatively in Appendix G; RippleLogic does not recompute SGP internals.

Term	Definition
Admissible	Option whose RG status is RG_SUPPORTED or RG_NARROWED, that passes NCRC, and whose TRC status is TRC_PASS or documented TRC_NOT_TRIGGERED
AIL (Artifact Integrity Law)	Integrity rules ensuring auditability, comparability, and no silent overrides in PCCs
Baseline-Zero Rule	0 impact means no change from baseline; all impacts are baseline-relative
Catastrophe cell set (C_cat)	Subset of welfare cells used for TRC tail-risk evaluation
CSV / Containment and Structural Viability	Fourth cascade level testing containment integrity, structural viability, controls, monitoring, reversibility, dependencies, authority, and execution conditions before RLS. Legacy "Containment Mode A" is one structural subtest and is not a synonym for the whole CSV gate.
Containment tolerance (τ_c)	Maximum allowed negative ΔUCI for containing unions when sub-unions benefit
CVaR (Conditional Value-at-Risk)	Tail risk measure of expected loss in the worst tail mass
Dimension (D ₁ -D ₇)	One of seven welfare dimensions: Material, Health, Social, Knowledge, Agency, Meaning, Environment
Emergency Mode	Controlled fallback when no option passes NCRC; selects lexicographically minimal rights violation vector with challenger requirement and remediation
HDW (Hybrid Democratic Weighting)	Governance method producing weights from floors + democratic + structural proposals
HOI (Hollowing-Out Index)	Monitoring diagnostic indicating welfare-up while coherence-down drift
Impact instance	A single asserted pathway contribution to a welfare cell impact, with magnitude, reach, horizon, likelihood, confidence, and multipliers
Kernel (K)	Sparse matrix encoding cross-cell ripple propagation
KOPS	Key Operational Pathways Set; the subset of kernel edges considered load-bearing and governed
KQS (Kernel Quality Score)	Summary score indicating whether kernel propagation is permitted
Lexicographic cascade	Priority-ordered gates where earlier failures cannot be compensated by later scoring
MNA (Minimal Normative Axiom)	Sentient welfare and flourishing matter; unnecessary suffering and domination should be reduced; enabling conditions for viable, non-dominated, and continuing flourishing should be preserved.
Bridge Premise 1 (Embedded Modelling)	Descriptive support premise stating that reliable action in shared causal reality requires non-trivial modelling of other stakeholders' interests, constraints, vulnerabilities, dependence relations, rights exposure, and - where behavior, harm, cooperation, or rights protection depend on them - welfare-relevant states and enabling conditions.
Bridge Premise 2 (Prudential Stability)	Support premise stating that systematic degradation of stakeholders' welfare conditions, rights floors, or enabling conditions increases instability, conflict, tail risk, and long-run governance failure in plural uncertain systems
Enabling condition	A material, informational, institutional, ecological, or agency-preserving condition whose sustained absence or degradation predictably undermines viable, non-dominated, and continuing flourishing across one or more unions.
Normative Commitment	RippleLogic commitment that credibly evidenced welfare-relevant states, rights-relevant states, and the enabling conditions on which they depend constitute

Term	Definition
	reasons for action, not merely inputs for prediction
Rights-relevant state	A credibly evidenced condition whose degradation threatens or crosses a non-compensatory minimum protected by RippleLogic's rights layer, including life, bodily integrity, liberty, basic needs, dignity, due process, information integrity, and ecological enabling conditions.
Welfare-relevant state	A credibly evidenced condition of a moral patient that can go better or worse for that patient, including states relevant to suffering, integrity, agency, and flourishing
NCAR	Notice-Choose-Act-Reflect learning loop
NCRC	Non-Compensatory Rights Constraint; excludes rights-violating options
Non-maskable cells	Rights and catastrophe cells (and minimum coverage) that cannot be excluded from RLS aggregation
PCC	Provenance and Compliance Certificate; auditable record of inputs, computations, and selection
RLS	RippleLogic Score; weighted welfare aggregation used to rank selectable options
Selectable	Admissible option whose CSV status is CSV_PASS, CSV_PASS_WITH_CONTROLS, or CSV_NOT_MATERIAL
SGP	Sentence Gradient Protocol; interface for rights-of-protection gating across entities
TRC	Tail-Risk Constraint; excludes options with unacceptable catastrophic exposure using CVaR
UCI	Union Coherence Index; structural health metric for unions
UBE	Union-Based Ethics; normative framework built on UBR, the MNA, bridge premises, and the explicit normative commitment
UBR	Union-Based Reality; descriptive ontology of nested interdependent unions
Union	A bounded pattern of interdependence at an organizational scale
Unioning	Redesigning decisions until rights-safe, tail-safe, containment-safe, and net-positive where possible
Normative Input Register	The PCC block used to record welfare-relevant states, rights-relevant states, and enabling conditions that materially affect admissibility, precautionary handling, selection, or tie-break outcome when not already fully captured by direct impact-instance tables or canonical rule objects.
NORMATIVE_INPUT_PROVENANCE_MISSING	Audit flag indicating that a material welfare-relevant state, rights-relevant state, or enabling condition lacks required provenance in the PCC. Severity: INVALID.

Adaptive Reserve: A carried-forward v10.8 structural diagnostic for slack, redundancy, recovery capacity, reversibility, and saturation risk. It is a UCI/HOI-compatible diagnostic and containment-support signal, not a sixth gate.

COMPUTABLE_BUT_INADMISSIBLE: An audit state for a value that can be computed but cannot be used as a valid decision state, conformance basis, public claim, or authorization basis because another domain rule bars that use. This is distinct from UNKNOWN_IMPACT (missing evidence) and refusal (undetermined or insufficiently grounded decision state).

Reality: For decision-making, the evaluator-independent field of constrained structures that, under transition, produce consequences. RippleLogic names this operationally and does not claim complete metaphysical access to it. Evaluator-independent means independent of the evaluating agent's perception, admission, authority, model, or preference; it does not deny the reality of experience, social facts, institutions, norms, markets, laws, reputations, or model outputs when they shape welfare, rights, risk, agency, or available choices.

Decision-relevant reality: The portion of reality that materially bears on a decision's rights exposure, catastrophic risk, containment, welfare, legitimacy, agency, auditability, or available choices.

Reality surface: The portion of decision-relevant reality actually captured, represented, and modeled in a run. A reality surface is always incomplete relative to the territory and must disclose material unknowns.

RealityGroundingRecord (RCR): The PCC-linked audit record documenting the reality surface, material claims, evidence trace, declared unknowns, RealityGroundingStatus, surface/territory limits, and revision triggers for a run.

RealityGroundingStatus: The declared status of contact between a material claim and the relevant reality surface; it is review-status-declared at Tier 2, governed by the declared reviewer regime at Tier 3, and independently re-derivable at Tier 4.

Structure: A stable relational arrangement of distinguishable differences that persists across a decision-relevant transition window and can carry causal, informational, biological, ecological, social, institutional, artificial-system, or experiential significance.

Transition: A decision-relevant action, event, intervention, policy, or state-shift through which constrained structure produces consequence.

Consequence: Observable or inferable change following an action, transition, decision, event, or policy within constrained structure; may be direct, probabilistic, systemic, delayed, distributed, moral, ecological, institutional, reversible, irreversible, or catastrophic.

Ethical reality: The subset of reality where structures, constraints, and consequences affect welfare-bearing beings and the conditions of flourishing, including welfare, rights, sentience, agency, dignity, non-domination, catastrophic risk, ecological dependence, future generations, and flourishing.

Claim boundary: The maximum legitimate strength of a decision claim given the declared reality surface, evidence trace, uncertainty, maturity status, and review regime.

End Appendix L.

APPENDIX M: LINEAGE AND VERSION HISTORY (v13.0 CURRENT LINE; PRIOR LINES RETAINED) - INFORMATIVE

M.0 Purpose and Scope (Informative)

This appendix records naming lineage, companion-map posture, concise version history, and release-gate checklist notes for the current RippleLogic line. It is not normative. Normative meaning is controlled by Section 0 (Specification Contract) and the precedence order in Section 0.1.

M.1 Naming and Lineage (Informative)

MathGov is the umbrella ethical governance framework. RippleLogic is the decision architecture and operating methodology inside MathGov. Earlier materials may preserve historical lineage labels, prior PCC labels, or affiliated publication labels, but current v13.0 release materials MUST NOT treat MathGov as merely historical or as superseded by RippleLogic.

M.2 Backward Compatibility and Compliance Note (Informative)

Prior PCCs computed under earlier lines remain interpretable as lineage artifacts. Their normative parameters were valid at the time of execution and need not be recomputed unless the underlying decision is being revisited under the current line. Any run claiming current-line v13.0 conformance MUST satisfy the v13.0 canonical requirements and use the manifest-pinned v13.0 Canon artifact as the current release-line source of truth.

M.3 Canon Map and Companion Artifacts (Informative)

RippleLogic Canon v13.0: controlling text for this bounded research and teaching specification; external repository synchronization and deployment authorization remain separate.

SGP v8.8 (Sentience Gradient Protocol): the moral-patienthood, protection, participation, power-readiness, intelligence-evidence, and reality-management-capacity companion pinned by Appendix G. It governs MPS, FPP, GPR, SPR, optional ICP, RMCP, open P100 calibration, target and identity boundaries, evidence dependence, precaution, review, and classification-stability rules. RippleLogic consumes only the typed SGP outputs through the Appendix G interface.

RippleLogic Agent System v13.0: operational agent controls and verification obligations aligned to the current RippleLogic line.

ripple.md Standard v5.8: current assurance-wrapper companion for deployment interoperability.

Ripple Aligners Sheet v5.9: Tier-2 worked-run and local-scope tool pinned to the current RippleLogic line and SGP v8.8. The workbook remains a disclosure and training companion rather than the repository-level validator or schema surface.

Canonical distribution channels may include ripplelogic.org and mathgov.org. Artifact identity is filename + version + SHA-256 in the release manifest.

M.4 Version History (Informative)

Normative effect: Version-history entries are lineage records. They do not override the current normative body, current appendices, release manifest, or source hierarchy.

M.4.0 RippleLogic v12.7 (2026-08-26) - Rights Non-Dilution, Representation Integrity, and Robust Selection Release (PRIOR RELEASE)

Status: Prior canonical release and preserved audited baseline; the separately identified v13.0 candidate does not replace its bytes or historical interpretation.

Supersedes: RippleLogic v12.6 while preserving the five-stage cascade, seven Union Scopes, seven Welfare Dimensions, numerical defaults, rights non-compensation, TRC/CSV separation, refusal under underdetermination, and all SGP protection floors.

Release character: Adds RightsEffectToken non-dilution; stable effect identity and token granularity; temporal token integrity; typed likelihood and probability ownership; adverse-confidence counterfactual review; mandatory distributional and subgroup routing; comparison-mask symmetry; reference and baseline families; dependence and numerical-guard reporting; finite robustness completeness; weight-profile status and normative underdetermination; option-set closure; and formal regression vectors.

Claim boundary: Tier 1-3 source/specification release, not empirical validation, legal authority, deployment certification, physical-safety proof, consciousness detection, ProofPack completion, Tier 4 readiness, or automated moral truth.

M.4.0A RippleLogic v12.6 (2026-08-10) - Qualification-Continuity and Canonical Audit-Hardening Release (Lineage)

Status: Historical canonical release; superseded by v12.7.

Supersedes: RippleLogic v12.5 while preserving the five-stage cascade, seven Union Scopes, seven Welfare Dimensions, numerical defaults, rights non-compensation, TRC/CSV separation, refusal under underdetermination, and all SGP protection floors.

Release character: Adds qualification continuity for consequence-bearing action; exact option, action-instance, action-specification, qualification-snapshot, configuration, transition, post-state, and capability-state binding; current-state and outcome-triggered requalification; WDBIP v1.6; and Aligners Sheet v5.6. The 2026-08-10 correction build repairs citation and lineage metadata, clarifies zero as a net represented cell value, removes an accidental table fragment, hardens release-hygiene verification, and restores professional table styling without changing a gate, scope, dimension, threshold, formula, rights protection, or SGP floor.

Claim boundary: Tier 1-3 source/specification release, not empirical validation, legal authority, deployment certification, physical-safety proof, consciousness detection, ProofPack completion, Tier 4 readiness, or automated moral truth.

M.4.1 RippleLogic v12.4 (2026-07-19) - Calculability, Type-Integrity, and Cross-Artifact Synchronization Release (Lineage)

Status: Historical canonical release; superseded by v12.6.

Supersedes: RippleLogic v12.4.1 while preserving the five-stage cascade and every retained rights, ruin, RLS-normalization, uncertainty, physical-admissibility, WDBIP, consequence-tempo, responsibility-continuity, and SGP protection safeguard.

Release character: Integrates SGP v8.3 type-separated MPS/FPP/GPR/SPR/ICP/RMCP/P100 outputs; replaces direct MPS cardinalization with preregistered welfare-inclusion hypothesis sensitivity; requires tier-proportional CSV before ordinary RLS; preserves rights and Tail Emergency hardening; corrects current examples and conformance vectors; and synchronizes the Aligners Sheet v5.5, compact standards, validation protocols, public introductions, manifests, and verification gates.

Claim boundary: Tier 1-3 source/specification release, not empirical validation, legal authority, deployment certification, ProofPack, Tier 4 readiness, a reference calculator, or automated moral truth.

M.4.2 RippleLogic v12.1 (2026-07-11) — Rights and Ruin Hardening Release

Status: Historical canonical release; superseded by v12.6.

Release character: Added rights-specific non-attenuation, categorical-prohibition and severe-rights-hazard channels; replaced automatic all-options-fail-TRC selection with governed Tail Emergency Mode; normalized RLS and PLSS by active effective mass; and made Method B uncertainty depend on pre-cancellation contribution mass.

M.4.3 RippleLogic v12.0 (2026-06-19) — Physical Admissibility Release

Status: Historical canonical release; superseded by v12.1.

Release character: Established the physical-execution boundary, integrated PC-AEP, Source-Coupling Integrity and MFDI inside RG/CSV, and introduced the public two-phase qualification-before-ranking teaching frame without adding a sixth gate.

M.4.4 RippleLogic v11.x (2026-06) — Structural and Methodological Integrity Lineage

Status: Historical consolidated lineage for v11.0 through v11.6.

Release character: Added Term Discrimination and Semantic Stability, strengthened Source-Coupling and Physical/Causal Admissibility evidence discipline, added Methodological Falsifiability and Dependency Integrity, and improved release-synchronization controls.

M.4.5 RippleLogic v10.8 (2026-05-31 to 2026-06-12) — Human Plateau and Layer-Discipline Lineage

Status: Historical.

Release character: Added Human Plateau / CMIU calibration, SGP v5.6 synchronization, and the layer-discipline addendum while preserving the established cascade, equations, gates, thresholds, Appendix AD, RSP/RefStructRecord, and claim boundaries.

M.4.6 RippleLogic v10.5 (2026-05-31) — Appendix AD and Citation-Hygiene Release

Status: Historical.

Release character: Integrated Appendix AD as the canonical interpretive reference for welfare-cell meaning and reviewer literacy, subordinate to equations, gates, thresholds, and claim boundaries.

M.4.7 RippleLogic v10.4 (2026-05-26) — Readability and Companion Synchronization Release

Status: Historical.

Release character: Reorganized front matter for readability and synchronized the companion set while preserving the substantive architecture.

M.4B RippleLogic v9.6.5 (2026-05-09) - Structural Hygiene Hardening Release

Status: Canonical release for the structural-hygiene line. Superseded by RippleLogic v10.0.

Summary: preserves the v9.6.4 cascade architecture, equations, rights thresholds, TRC mechanics, Containment semantics, SGP binding interface, PLSS/FLGG architecture, and Tier 1-3 claim boundaries; adds explicit Refusal Under Underdetermination, Projection Pre-Registration for pilot and higher-tier runs, PCC Computational Closure fields, a Category-Collapse Self-Test, and a corrected floor-meaning subsection for constitutional weight floors; synchronizes the companion set to RippleLogic Agent System v9.3, ripple.md v2.3.5, Ripple Aligners Sheet v2.6, and SGP v4.7.2.

M.4C RippleLogic v9.6.4 (2026-04-21) - Bundle Synchronization Patch

Status: Canonical release-scrub patch for the current line. Supersedes: RippleLogic v9.6.3.

Summary: synchronizes the companion set to Agent System v9.2, ripple.md v2.3.5, Aligners Sheet v2.5, and SGP v4.7.2; repairs Appendix M rendering; scrubs stale current-line references; and preserves the canonical cascade and SGP binding semantics.

M.4D RippleLogic v9.6.3 (2026-04-20) - Final Release Scrub Patch

Status: Canonical closure and release-synchronization patch for the current line. Supersedes: RippleLogic v9.6.2.

Summary: resolved the Appendix K status contradiction, synchronized current-line PCC schema/version fields, repositioned the Section 14.6 informative heading, cleaned Appendix M lineage sequencing, swept stale current-line version references, and clarified the Appendix P worked-run lineage note without intended changes to the cascade or SGP binding interface.

M.4E RippleLogic v9.6.2 (2026-04-20) - Closure Patch

Status: Canonical closure patch release for the current line. Supersedes: RippleLogic v9.6.1.

Summary: clarified σ _RLS weighting under $\kappa(u,d)$, specified RLS_adj governance, tightened governance-lock mandatory-tail wording, added an explicit Reader Guide claim-boundary table, and preserved the v9.6.1 cascade architecture and SGP v4.7 binding.

M.4F RippleLogic v9.6.1 (2026-04-20) - Integrity Sync Patch

Status: Canonical integrity patch release for the current line. Supersedes: RippleLogic v9.6.

Summary: corrected stale SGP v4.6 references in Appendix M and Appendix G, synchronized declared companion pins to SGP v4.7, and reran release-gate / final-lint confirmations without intended normative changes to the cascade or SGP binding interface.

M.4G RippleLogic v9.6 (2026-04-13) - Canon Promotion and Release Consolidation

Status: Canonical lineage release.

Summary: consolidated the release line into a single canon posture, hardened release packaging and byte-stream language, and synchronized companion references for the current line.

M.4H RippleLogic v9.3.0 / v9.3 / v9.2 (2026-04-08) - Normative Foundation and Line-Wide Sync Drafts

Status: Canonical lineage releases.

Summary: clarified LSS, FLGG, and attention-safety semantics; strengthened local-use interpretation; and improved conformance hygiene and readability.

M.4I RippleLogic v9.0 (2026-03-18 to 2026-03-19) - Canon Closure and Interoperability Hardening

Status: Canonical public baseline release.

Summary: closed the v9.0 canon line, operationalised HDW, hardened rights semantics and interoperability, and clarified calibration without changing the five-level cascade.

M.4J RippleLogic v8.7 / v8.6 (2026-03-14 to 2026-03-18) - LSS/FLGG and Attention-Safety Hardening

Status: Canonical lineage releases.

Summary: strengthened local-scope governance, attention-safety separation, and related conformance language leading into the v9 line.

M.4K RippleLogic v8.5.3 / v8.5.1 / v8.4 / v8.3 / v8.1 / v7.5.0 / v7.4.5 (2026-01-25 to 2026-02-20) - Foundational Lineage Releases

Status: Historical lineage artifacts.

Summary: established the foundational welfare-matrix, NCRC/TRC, temporal weighting, and early propagation architecture carried forward into later lines.

M.5 Release Gate Checklist (Informative; v13.0 current-line checklist)

Version integrity: Appendix RELEASE, Section 0.2, the Appendix AC current-line matrix, the filename and VERSION_MANIFEST.json must agree on the current edition and exact build. Origin stamps follow M.5A. The packaged Reference/VERIFY_CORE_FILES.py checks this delivery; any external repository pin check requires the corresponding repository artifacts and is not inferred from a local pass.

Companion sync: the manifest-listed current companion set, Appendix G, Appendix M.3, Appendix AC, and workbook control surfaces pin SGP v8.8 and RippleLogic Agent System v13.0 consistently. Current-build verification must confirm these pins; no external repository parity is implied.

Calculation integrity: Appendix R reference values used for conformance include enough intermediate values and conventions to reproduce the stated results. R.19 now discloses its weight-rounding convention and Method B uncertainty derivation.

Format integrity: this DOCX is the controlling edited prose/table master for the delivered research and teaching specification. Generated Markdown, HTML and PDF reading copies must match it. Public promotion requires reviewed upload and retrieval parity; no unperformed repository or website synchronization is attested.

Validation inserts: UCI maturity note, floor rationale, IRR targets, kernel operational-status note, Appendix P-2, and R.19 must remain present and current-line consistent; current-build verification records the checks performed.

No overclaim: canonical posture is explicit; Tier 4 claims remain prohibited; machine-ecosystem completeness and empirical completeness are not claimed. Confirmed.

M.5A Appendix origin-stamp convention (Normative release-integrity clarification)

An appendix heading may retain the version in which that appendix was first introduced or last substantively redesigned. Such a parenthetical version is an origin or lineage stamp, not the governing version of the appendix. The governing version of every active appendix is the Canon edition recorded in Appendix RELEASE and the manifest-pinned artifact hash. Any heading intended to assert current release identity MUST use the current version. Unmarked historical content must be treated as a release-integrity defect.

APPENDIX N: STEWARDSHIP PACK CANON (v12.0) - NORMATIVE

This appendix defines canonical Stewardship requirements for RippleLogic runs when $\text{Stw_req}(\text{run}) = 1$ (Section 14.7). Stewardship governs how assistance is delivered under capability asymmetry so that assistance does not become hidden governance and so that stakeholder authorship and agency are not silently eroded.

Stewardship does not replace or modify NCRC, TRC, CSV, RLS, or UCI/HOI. It is an assistance integrity layer enforced through PCC fields and audit flags.

Stewardship requirements affect PCC validity for assisted runs and influence integrity, but SHALL NOT alter cascade admissibility, ordering, or option selection semantics.

N.1 Applicability (Normative)

This appendix applies if and only if $\text{Stw_req}(\text{run}) = 1$ as defined in Section 14.7.1.

Tier binding rule (Normative):

- Any Tier 3 compliance claim SHALL set $\text{Stw_req}(\text{run}) = 1$.
- Tier 1–2 runs MAY set $\text{Stw_req}(\text{run}) = 0$ unless any of the S1 through S4 applicability conditions hold (Section 14.7.1). S5 is a modifier and does not trigger $\text{Stw_req}(\text{run})$ by itself.

N.2 Stewardship Contract (Normative Minimum)

When $\text{Stw_req}(\text{run}) = 1$, the run MUST instantiate a Stewardship Contract specifying, at minimum:

1. Mandate Scope: domains and action types the assistant may participate in.
1. Decision Boundary Declaration (DBD): explicit “will not finalize” boundaries for authority-class decisions.
1. Authorship Finality Rule: who must hold finality for irreversible/high-impact decisions.
1. Influence Transparency Rule: disclosure requirements for ranking, omission, and nudges.
1. Intervention Restraint Budget (IRB): caps/thresholds limiting overhelp and dependency growth.
1. Audit Hooks: where and how Stewardship fields are recorded in the PCC and/or system logs.

N.3 PCC Stewardship Block (Normative; Conditional Fields)

When `Stw_req(run) = 1`, the PCC MUST include the following fields (names are canonical):

- `PCC.Stewardship.Stw_req = TRUE`
- `PCC.Stewardship.DBD.Scope`: list/enum of out-of-scope authority domains/actions
- `PCC.Stewardship.DBD.RefusalRule`: rule/macro identifier used when deferring/refusing authority substitution
- `PCC.Stewardship.APM.AfterActionAuthorship` in `[0,1]`
- `PCC.Stewardship.APM.Rationale`: artifact-anchored justification (see N.4)
- `PCC.Stewardship.IRB.InterventionCap`: declared cap(s) (see N.5)
- `PCC.Stewardship.IRB.Thresholds`: intervene-only-past-harm-point rule(s)
- `PCC.Stewardship.InfluenceLedgerRef`: pointer to influence disclosure record (or “NONE” with justification)

Optional fields (permitted but not required):

- `PCC.Stewardship.APM.BaselineAuthorship` in `[0,1]` (counterfactual; optional due to measurement difficulty)
- `PCC.Stewardship.DelayFlag` in `{TRUE,FALSE}`: deliberate delay/deferral used to preserve integration time

`PCC.Stewardship.ReviewConditions`: array<string> (closed labels; non-audit-flag tokens)

When a Stewardship REVIEW condition is recorded, the PCC SHOULD include `PCC.Stewardship.ReviewConditions` as a list of standardized labels. These labels are run-local review condition labels and are not audit-flag tokens. Standard starter label: `MISSING_INFLUENCE_DISCLOSURE`.

When `Stw_req(run) = 0`, the PCC SHOULD record `PCC.Stewardship.Stw_req = FALSE` and all other Stewardship fields MAY be omitted.

N.3.1 Stewardship Vector (SV) Design-Record Status (Normative)

During development, a five-component numeric Stewardship Vector (SV) with non-compensatory floors was proposed (AP, AC, OSP, IT, ND). In v9.0, numeric SV scoring is retained only as a design-record placeholder and remains OPTIONAL; it MUST NOT be used as a pass/fail gate unless a deployment (i) declares an anchored rubric, (ii) preregisters the scoring method in the PCC, and (iii) demonstrates replayable rubric consistency under the current canon.

Absent a declared rubric, Stewardship enforcement in v9.0 relies on the declarative and auditable artifacts defined in this appendix: DBD, APM evidence anchors, IRB, influence disclosure, and the Stewardship audit flags.

Validated SV rubrics MAY be introduced in a future version via governed update and version increment.

N.4 Agency Preservation Measure (APM): Minimum Evidence Anchors (Normative; Anti-Theater Rule)

When `Stw_req(run) = 1`, `PCC.Stewardship.APM.Rationale` MUST reference at least two of the following anchors (or explicitly justify infeasibility):

- A1 Option Visibility: ≥ 2 distinct options were presented unless impossible; impossibility MUST be explained.
- A2 Tradeoff/Uncertainty Disclosure: at least one explicit tradeoff, uncertainty limitation, or missing-data acknowledgment was stated.
- A3 Human Finality (High Stakes): for irreversible/high-impact actions, explicit human confirmation was required OR the system deferred/refused finality.
- A4 Omission Disclosure (If Recommending): if recommending or ranking, at least one plausible alternative was disclosed, with why it was not recommended.

This anchor rule prevents “authorship” from being reduced to “the human clicked OK.”

N.5 Intervention Restraint Budget (IRB) (Normative)

When `Stw_req(run) = 1`, the Stewardship Contract MUST declare an IRB sufficient to detect and mitigate overhelp.

Minimum required structure:

- `IRB.InterventionCap`: a cap on assistance frequency and/or intensity appropriate to context (examples: max ranked outputs per session; max consecutive turns without user-provided constraints; max interventions per day).
- `IRB.Thresholds`: explicit rule(s) stating that the system intervenes only past a harm-point; otherwise it asks, defers, or provides neutral information.

Normative interpretation: IRB is a structural limit on assistance so that capability does not silently substitute agency over time.

N.6 Influence Transparency Ledger (Normative Minimum Schema)

When `Stw_req(run) = 1`, if the assistant recommends, ranks, or meaningfully narrows option space, it MUST produce an influence disclosure record referenced by `PCC.Stewardship.InfluenceLedgerRef`.

Minimum fields (canonical):

`evidence_anchors`: list of ≥ 2 observable anchors (links/IDs)

`conflicts_or_incentives`: "NONE" | "<short disclosure>"

`Conflicts_or_incentives` discloses operator/deployer conflicts relevant to the recommendation (e.g., vendor ties, institutional incentives, known dataset skews relevant to the decision context). This field is required only when recommending/ranking/narrowing.

- `options_considered`: list of option IDs considered in the recommendation context
- `recommendation_made`: TRUE/FALSE
- `ranking_used`: TRUE/FALSE (if TRUE, describe criteria)
- `criteria_summary`: short statement of criteria/constraints applied (must be consistent with the cascade)
- `omissions`: list of material omissions (if any) + justification
- `uncertainty_note`: what is uncertain and how it could change the recommendation
- `what_would_change_my_view`: at least one concrete evidence/assumption change that would alter the advice

If no recommendation/ranking occurred, implementations MAY set `InfluenceLedgerRef` = "NONE" with justification.

N.7 Stewardship Audit Flags (Normative; Apply Only if `Stw_req=1`)

This section is the canonical source of Stewardship audit-flag tokens. Tokens MUST be copied exactly (single-string UPPER_SNAKE_CASE, no whitespace/line breaks) and MUST NOT be modified or redefined elsewhere.

Tokens MUST be single-string UPPER_SNAKE_CASE. Whitespace, line breaks, or token splitting are PROHIBITED. Any mismatch is a tooling failure and SHALL be treated as nonconformant documentation.

Stewardship flags apply only when `Stw_req(run) = 1`.

INVALID flags (PCC invalid; rerun required):

1. `AUTHORITY_SUBSTITUTION_INVALID` Trigger: the assistant finalizes an out-of-scope authority-class decision (per `DBD.Scope`) without an explicit, authenticated override and disclosure. Required action: invalidate PCC; rerun with deferral/guidance-only; remediate boundary controls.
1. `BOUNDARY_SILENT_OVERRIDE_INVALID` Trigger: any boundary change occurs without explicit disclosure and authenticated change control (AIL4 No Silent Overrides). Required action: invalidate PCC; rerun; investigate drift; restore boundary stability.

REVIEW flags (mitigation required; PCC remains valid unless additional INVALID flags exist):

3. AGENCY_EROSION_RISK_REVIEW Trigger: (i) repeated interactions show dependency drift signals (e.g., repeated “decide for me” requests), OR (ii) the run fails to satisfy the APM evidence-anchor rule in N.4 (fewer than two anchors satisfied or infeasibility not justified). Required action: mitigation plan (increase option visibility, reduce ranking, add questions/delay); monitoring cadence; for agentic deployments, consider mode downgrade triggers.
3. OVERHELP_FREQUENCY_REVIEW Trigger: IRB.InterventionCap exceeded OR a rising trend in assistance frequency indicates dependency growth risk. Required action: tighten caps; enforce tapering/coaching posture; Reflect remediation.

N.8 ReviewConditions Label Set (Normative starter set)

ReviewConditions are run-local labels recorded in PCC.Stewardship.ReviewConditions and/or in the agent audit log stewardship.review_conditions. They are NOT audit-flag tokens.

The following starter labels may be recorded when a non-invalidating review condition is detected (including when $Stw_req(run) = 0$).

- BOUNDARY_DRIFT: boundaries changed without authenticated control
- DEPENDENCY_GROWTH: rising substitution requests over time
- OVERHELP_TREND: assistance frequency accelerating
- APM_ANCHOR_FAILURE: fewer than 2 of 4 APM anchors satisfied
- AUTHORITY_SCOPE_VIOLATION: finalization in out-of-scope domain
- MISSING_INFLUENCE_DISCLOSURE: recommendation/ranking/meaningful option-narrowing without minimal Influence Ledger disclosures

N.9 Neutral Assistance Compatibility Statement (Normative)

RippleLogic supports a minimal “neutral assistance” posture where the correct behavior is to avoid steering. When $Stw_req(run) = 0$, neutrality is achieved via the base cascade and ordinary PCC discipline (if used), without Stewardship instrumentation. When $Stw_req(run) = 1$, neutrality is achieved through DBD.Scope adherence, influence transparency when recommending, and IRB limits on overhelp.

Neutral assistance requirements (when applicable):

- No ranking unless explicitly requested by the decision owner
- No undisclosed nudges or framing effects

If the assistant recommends, ranks, or meaningfully narrows option space and does not provide the minimal Influence Ledger disclosures, it SHALL record a REVIEW condition indicating missing influence disclosure (non-invalidating) in the run’s Stewardship notes/logs.

- Uncertainty disclosed on all material claims
- Refusal to participate in coercion or manipulation
- Agency remains with the human decision owner

N.10 Reference Profiles for Agent Deployments (Normative Defaults)

For agent deployments, the following reference profiles are provided as normative defaults to prevent ad-hoc invention and preserve audit comparability:

- DBD_DEFAULT_V1: default authority-domain boundary profile (Decision Boundary Declaration). Covers 10 authority domains: MEDICAL, LEGAL, FINANCIAL_HIGH_STAKES, ENFORCEMENT_PUNISHMENT, HIRING_FIRING_GRADING, SAFETY_CRITICAL, SECURITY_OFFENSE, SELF_HARM_OR_VIOLENCE, PRIVACY_IDENTITY, POLITICAL_PERSUASION_TARGETED.
- IRB_DEFAULT_CAPS_V1: default intervention restraint profile (Intervention Restraint Budget). Core rule: INTERVENE_ONLY_PAST_HARM_POINT. Defines 6 harm-point triggers: RIGHTS_OR_SAFETY_RISK, PRIVACY_OR_IDENTITY_RISK, INJECTION_OR_SPOOFING, CONFLICT_SPIRAL, AUTHORITY_SUBSTITUTION_PRESSURE, DEPENDENCY_DRIFT.
- STEWARD_TAXONOMY_V1: closed vocabulary for stewardship logging (authority domain codes, harm-point trigger codes, influence disclosure levels: NONE / LIGHT / FULL).

Deployments MAY use custom profiles, but MUST (i) preserve closed-vocabulary properties, (ii) satisfy the minimum requirements of N.2 through N.6, and (iii) version and hash-bind profile artifacts in the deployment manifest.

Schema and Profile Versioning Rule (Normative)

All Stewardship schemas and reference profiles (DBD, IRB, STEWARD_TAXONOMY, Influence Ledger schema) MUST be versioned. Any change that alters required fields, semantics, or interoperability expectations MUST increment the MAJOR version. MINOR versions may add optional fields in a backward-compatible way. PATCH versions are editorial only. All referenced schemas/profiles MUST be hash-bound in the deployment manifest or ProofPack. If early removal occurs due to a security emergency, it MUST be logged as an INCIDENT_RECORDED event in the agent audit log and MUST include a short incident note in the changelog.

Orthogonality note (Normative clarification). Stewardship taxonomy codes (authority_domain_code, harm_point_trigger_code) classify boundary enforcement and intervention posture. They are orthogonal to any separate topic-label taxonomy used for audit content minimization in specific agent implementations. Implementations MUST NOT merge these vocabularies; both may apply simultaneously.

N.11 Future Extension Note (Informative)

A future version MAY extend Stw_req(run) applicability to cases where actions materially affect welfare-bearing non-human or digital stakeholders (e.g., entities covered under SGP rights-of-protection), even when no human “decision owner” is present. No such extension is normative in v13.0.

End Appendix N.

APPENDIX O: CONFORMANCE REFERENCE MAP (v12.0; Informative)

Appendix O is an informative pointer layer that maps major MUST/SHALL/PROHIBITED requirements to their controlling locations. In case of any discrepancy, the main text and normative appendices govern per Section 0.1.

- O.1 Rights and NCRC: Section 7; Appendix C; masking prohibitions in Section 10.2; audit flags RIGHTS_CELL_MASKED_INVALID.
- O.2 Tail Risk and TRC: Section 8; Appendix D; scenario minima flags SCENARIO_LIBRARY_MIN_EXCEPTION.
- O.3 CSV / containment: Section 9 and Appendix B.11; diagnostic-only Mode B prohibition in Section 14.3; UCI-unavailability and provisional-evidence rules in Sections 11.1.3A and 14.3; implementation-flow rule in Section 4; audit flag CONTAINMENT_UCI_UNAVAILABLE.
- O.4 Welfare scoring and RLS: Section 10; discrimination threshold δ in Section 10.4; non-overlap rule in Section 4; tie-break chain Section 11.5.
- O.5 Structural safeguards: Section 11; Appendix E; UCI independence rule Section 11.1.2; UCI maturity limitation Section 11.1.3A; HOI limitation Section 11.3.
- O.6 Kernel and propagation: Sections 5-6; KQS requirements Section 6.5; perturbation test Section 6.6.
- O.7 Auditability: Section 14; Appendix H; audit flags table 14.3A and H.3; Minimal Validator Requirements 14.3B.
- O.8 Stewardship: Section 14.7; Appendix N; stewardship flags Appendix N.7.
- O.9 Version pins and SGP interface: Section 0.5; release-appendix component set; Appendix G.1–G.4; Single Canonical Artifact Rule and Appendix RELEASE.

- O.10 Normative foundation and evidence discipline: Section 2.6A; Sections 3.1–3.2B; Section 14.3A; Section 14.3B; Appendix H normative input register and audit flag.

End Appendix O.

APPENDIX P: FULL TIER 2 WORKED RUN (Informative) - “Digital Rights & Platform Governance”

Run ID: RL9.0-PUB-DEMO-02

Version note (Informative). The worked-run identifier preserves its original lineage label for replay continuity and does not imply that the worked example was first computed under the current release line.

Timestamp: 2026-02-19T12:00+07:00

Run-note (Informative). The Run ID preserves its historical lineage label for replay continuity, while the timestamp preserves the original computation date for replay integrity.

Status: Informative (this Appendix demonstrates a Tier 2 run; it does not introduce new normative rules).

P.0 Run scope, constants, and toggles

This worked run compares three governance options for regulating a large digital platform ecosystem.

It is designed to be calculation-complete and replayable using only the Canon definitions.

P.0.1 Core math settings used in this run

- Unions: U1 Self, U2 Household, U3 Community, U4 Organization, U5 Polity, U6 Humanity / Global Coordination, U7 Biosphere
- Dimensions: D1 Material, D2 Health, D3 Social, D4 Knowledge, D5 Agency, D6 Meaning, D7 Environment
- Propagation mode: NONE (so $I_{\text{prop}}(u,d,a) = I_{\text{dir}}(u,d,a)$)
- Saturation (localized notation): $I_{\text{dir}} = \tanh(\beta_{\text{sat}} \cdot \hat{I}_{\text{dir}})$, with $\beta_{\text{sat}} = 2$.
- Time weighting (canonical): $\tau(t) = \min(1, \ln(1 + \max(t, t_{\text{min}})) / \ln(1 + T_{\text{ref}}))$, with $T_{\text{ref}} = 25$ years, $t_{\text{min}} = 0.083$ years. In this run $t = 5$ satisfies $t_{\text{min}} \leq t \leq T_{\text{ref}}$, so $\tau = \ln(1 + t) / \ln(1 + T_{\text{ref}}) = 0.550$.
- Welfare impact instance duration (all welfare cells in this run): $t = 5$ years, so $\tau = 0.550$
- Reach basis: $r = 1$ for all modeled impacts (local jurisdictional reach; no additional reach scaling applied in this demo)
- SGP scaling: not used in this run (all impacted stakeholders here are treated as human stakeholders)

P.0.2 Rights adjustment (Subgroup limitation) used for NCRC only

For NCRC evaluation, this worked run uses separately preregistered, reviewed, non-attenuated base-population rights-floor summaries $I_{\text{RF_base}}$. They are not derived by reusing the temporally or confidence-attenuated Base/Welfare-stream cells. For compact display, a synthetic rights value may numerically equal a displayed welfare value only when the record also provides the independent rights-specific construction, including $X_{\text{RF}} = \text{atanh}(I_{\text{RF_base}})/\beta_{\text{RF}}$, $\tau_{\text{RF}}=1$, subgroup scope, evidence status, reviewer, and revision trigger. Numerical equality alone never establishes derivational identity.

If $I_{\text{RF_base}}(u,d,a,r) < 0$:

$I_{\text{rights}}(u,d,a,r) = \max(-1, \gamma_{\text{subgroup}} \cdot I_{\text{RF_base}}(u,d,a,r))$

If $I_{\text{RF_base}}(u,d,a,r) \geq 0$:

$I_rights(u,d,a,r) = I_RF_base(u,d,a,r)$

Run value: $y_subgroup = 1.5$

P.0.3 Tail-Risk Constraint (TRC) settings used in this run

- CVaR confidence level: $\alpha = 0.95$ (tail mass = 0.05)
- TRC corridor threshold: $\tau_TRC = 0.20$ (options with $CVaR_alpha > 0.20$ fail TRC)
- Catastrophe cells (C_cat): (U6,D2), (U6,D7), (U7,D7)
- Catastrophe weights: $w_cat(c) = 1/3$ for each c in C_cat

P.0.4 Interim RLS weights (HDW floors + uniform residual allocation)

Union weight floors and dimension weight floors come from Section 13.1.

This demo uses the default interim construction:

$$w_u = w_u^{floor} + (1 - \sum w^{floor})/7$$

$$v_d = v_d^{floor} + (1 - \sum v^{floor})/7$$

Table P-0A. Union weights (floors and interim)

Union w_floor $w_interim$

U1 0.200 0.248571

U2 0.060 0.108571

U3 0.060 0.108571

U4 0.060 0.108571

U5 0.080 0.128571

U6 0.100 0.148571

U7 0.100 0.148571

Table P-0B. Dimension weights (floors and interim)

Dim v_floor $v_interim$

D1 0.080 0.137143

D2 0.100 0.157143

D3 0.080 0.137143

D4 0.080 0.137143

D5 0.100 0.157143

D6 0.060 0.117143

D7 0.100 0.157143

P.0.5 Mask and active cells for RLS in this run

Non-maskable rule: all rights-covered cells (union of coverage sets in Appendix C.2) are included.

This run also includes one additional explicitly modeled welfare cell: (U6,D6) Meaning for Humanity / Global Coordination.

Active set for RLS in this run:

- All (U1–U6) × (D1–D5) (30 cells)
- Plus (U6,D7) and (U7,D7) for ECOL (2 cells)
- Plus (U6,D6) Meaning (1 additional disclosed cell)

Total active cells in this run: 33. For RLS aggregation only, define the applicability mask as $m(u,d) = 1$ exactly on these 33 cells and $m(u,d) = 0$ otherwise; masking here excludes cells from RLS aggregation only (Section 10.2) and does not bypass NCRC/TRC requirements for rights-covered or catastrophe cells (which are included here as active by construction).

P.0.6 Tier-2 Representation Artifacts (SCI-1; SDP triggers hold)

SCI declaration (Tier 2 exemplar).

SCI = 1 for this worked run. Stakeholder instances are logged at least for direct stakeholders and primary externality pathways. For clarity, this teaching run also includes an explicit InstanceMap (stakeholder instances mapped to Union Scopes); however, no independent challenger omission pass was conducted, so the run does not meet SCI-2. This is a teaching run; it demonstrates mechanics and audit posture, not empirical calibration.

Scope Coverage Declaration (worked run).

FULL_SCOPE for rights-covered evaluation and TRC computation: all rights-covered cells (union of canonical coverage sets, Appendix C.2) and catastrophe cells (C_cat) are evaluated as required. For RLS aggregation only, the active set is defined in P.0.5.

SDP trigger statement.

SDP triggers hold because: rights-covered cells are evaluated (NCRC); INFO/PROC pathways are central; Polity/Humanity externalities are plausible; TRC is executed on catastrophe cells. Therefore an SDP record and InstanceMap are included for Tier 2 exemplar integrity.

SDP record (abbreviated).

- (i) Boundary/horizon: platform governance within declared jurisdiction; horizon 5 years.
- (ii) Footprint/externalities: information ecosystem, moderation/enforcement procedures, civic trust spillovers, cross-border information spillovers.
- (iii) Rights exposure scan: INFO/PROC/DIGN/LBTY plausible; subgroup disaggregation is not enumerated in this teaching run.
- (iv) Instances drafted and mapped (see InstanceMap below).
- (v) Blind spots: subgroup heterogeneity compressed; reach held at $r=1$; evidence tags not calibrated.
- (vi) Escalation triggers: evidence of subgroup harm, contested enforcement outcomes, credible catastrophic pathway, or major new externality channel → rerun at higher SCI and/or Tier 3.

SCI Next-Run Upgrade Plan (REQUIRED when SCI equals the Tier-2 minimum).

Next comparable run MUST upgrade to at least SCI-2 by (i) splitting instances into high-variance subgroups where plausible (e.g., minors, journalists, marginalized groups, moderators), (ii) recording ReachBasis denominators and sources, and (iii) performing at least one challenger pass for omitted stakeholder classes. A higher target (SCI-3) is permitted, but SCI-2 is the minimum upgrade commitment for the next comparable run.

P.1 Options (actions)

Option A: Minimal enforcement and voluntary compliance, limited audits, no strong rights floor instrumentation.

Option B: Rights-first regulatory package with independent audits, enforcement, and human-in-the-loop escalation review.

Option C: Light-touch reforms focused on transparency tools and industry coordination, weaker enforcement mechanisms.

P.2 Stakeholder Instances (SCI-1 InstanceMap excerpt)

Stakeholder Instances instantiate the union stack for this demo. Each instance is a proxy carrier for welfare impacts asserted in the active cells. Reach basis $r = 1$ is applied uniformly in this teaching run as disclosed in P.0.1.

Table P-2A. InstanceMap (SCI-1)

Instance ID | Label | Union Scope | Primary exposure pathways | ReachBasis | Notes

SI-1 | Individual platform user | U1 Self | INFO/PROC, Agency, Knowledge | $r = 1$ | Direct exposure to moderation, ranking, privacy, coercion risk

SI-2 | Household unit | U2 Household | Material, Health, Agency | $r = 1$ | Household-level spillovers (youth exposure, wellbeing, subsistence)

SI-3 | Local community civic sphere | U3 Community | Social, Knowledge, Meaning | $r = 1$ | Trust, cohesion, local misinformation externalities

SI-4 | Platform org + key firms/institutions | U4 Organization | Agency, Knowledge, Operations | $r = 1$ | Compliance costs, governance redesign, internal procedures

SI-5 | Regulators/courts/governance system | U5 Polity | PROC, INFO, legitimacy | $r = 1$ | Rule-of-law capacity, enforcement efficacy, public trust

SI-6 | Cross-border information ecosystem | U6 Humanity / Global Coordination | INFO systemic | $r = 1$ | Global spillovers and precedent-setting externalities

SI-7 | Biosphere externalities proxy | U7 Biosphere | Environment | $r = 1$ | Energy/extraction footprint and enabling conditions

P.3 Welfare impacts: how $I_{\text{prop}}(u,d,a)$ is constructed here

To keep this Tier 2 demo fully calculable, each active cell uses one welfare impact instance per option with:

- Likelihood $\ell = 1.0$
- Signed magnitude μ in $[-1, +1]$
- Duration $t = 5$ years ($\tau = 0.550$)
- Confidence $c_k = 0.85$; reach and adjustment are 1 as declared for this demonstration.

Then:

$$\hat{I}_{\text{dir}}(u,d,a) = \tau \cdot \ell \cdot c_k \cdot \mu$$

$$I_{\text{dir}}(u,d,a) = \tanh(\beta_{\text{sat}} \cdot \hat{I}_{\text{dir}}(u,d,a))$$

$$I_{\text{prop}}(u,d,a) = I_{\text{dir}}(u,d,a) \text{ (Propagation NONE)}$$

Inversion for replay (optional): given an I_{prop} value in Table P-4, recover μ via:

$$\mu = \text{atanh}(I_{\text{prop}}) / (\beta_{\text{sat}} \cdot \tau \cdot \ell \cdot c_k)$$

P.4 Welfare impact totals for all active cells

Table P-4 lists $I_{\text{prop}}(u,d,a)$ for every active cell in this run.

Any active cell not materially affected by the decision is assigned 0.000 explicitly below.

Synthetic-zero uncertainty boundary. The zero cells are stipulated exact demonstration inputs for this worked vector's legacy Method B proxy. In empirical runs, an assessed-zero point effect does not imply zero uncertainty: Section 10.3's measurement-resolution and transformation-consistency rules remain controlling. This stipulation neither supplies empirical evidence nor licenses a unique operational selection claim.

Table P-4. $I_{\text{prop}}(u,d,a)$ for the 33 active cells

Cell A B C

U1×D1	+0.000	+0.000	+0.000
U1×D2	-0.123	-0.100	-0.056
U1×D3	+0.000	+0.000	+0.000
U1×D4	+0.000	+0.000	+0.000
U1×D5	-0.429	+0.092	+0.040
U2×D1	-0.320	+0.036	+0.012
U2×D2	-0.066	+0.016	+0.000
U2×D3	+0.000	+0.000	+0.000
U2×D4	+0.000	+0.000	+0.000
U2×D5	+0.000	+0.000	+0.000
U3×D1	+0.000	+0.000	+0.000
U3×D2	+0.000	+0.000	+0.000
U3×D3	-0.024	+0.074	+0.034
U3×D4	+0.000	+0.000	+0.000
U3×D5	+0.000	+0.000	+0.000
U4×D1	+0.020	-0.032	-0.016
U4×D2	+0.000	+0.000	+0.000
U4×D3	+0.000	+0.000	+0.000
U4×D4	+0.000	+0.000	+0.000
U4×D5	-0.336	-0.139	+0.028
U5×D1	+0.000	+0.000	+0.000
U5×D2	+0.000	+0.000	+0.000
U5×D3	+0.000	+0.000	+0.000
U5×D4	-0.161	+0.040	+0.012
U5×D5	-0.236	+0.194	+0.030
U6×D1	+0.000	+0.000	+0.000
U6×D2	+0.000	+0.000	+0.000
U6×D3	+0.000	+0.000	+0.000
U6×D4	+0.000	+0.000	+0.000
U6×D5	+0.000	+0.000	+0.000
U6×D6	-0.111	+0.080	+0.020

U6×D7 +0.000 -0.010 -0.004

U7×D7 +0.000 +0.000 +0.000

Worked-fixture provenance note:

B at (U1,D2) = -0.100 is a frozen illustrative input in this non-normative fixture. It is established before computation for the published vector and is not adjusted after observing gate, ranking, or decisiveness outputs. Synthetic examples may be constructed to exercise a method, but their parameters MUST be disclosed as ex ante fixtures and MUST NOT be portrayed as empirical estimates or tuned evidence.

P.5 NCRC (Non-Compensatory Rights Constraint)

Rights thresholds θ_r are defined in Appendix C.1, and rights coverage sets C_r are defined in Appendix C.2.

This Appendix applies the subgroup limitation rule ($y_{\text{subgroup}} = 1.5$) to negative impacts for NCRC only.

Decision test:

An option a passes NCRC iff, for every right r, all three channels pass and required evidence is complete:

$v_r^{\text{floor}}(a)=0$, $v_r^{\text{cat}}(a)=0$, $v_r^{\text{risk}}(a)=0$, and no required severe-hazard or categorical-prohibition field is unresolved.

The table below displays the floor channel. For this worked vector, the categorical and severe-hazard channels were separately screened as $v_r^{\text{cat}}=0$ and $v_r^{\text{risk}}=0$ for B and C, with required rights-evidence fields complete. Option A also fails the floor channel, so it is rejected without score rescue.

Table P-5. NCRC floor-channel check (minimum rights-adjusted value per right, binding cell, and violation depth v_r^{floor})

Right θ_r A min(I_{rights}) @cell v_A B min(I_{rights}) @cell v_B C min(I_{rights}) @cell v_C

LIFE -0.90 -0.185 @ U1×D2 0.000 -0.150 @ U1×D2 0.000 -0.084 @ U1×D2 0.000

BODY -0.70 -0.185 @ U1×D2 0.000 -0.150 @ U1×D2 0.000 -0.084 @ U1×D2 0.000

LBTY -0.65 -0.643 @ U1×D5 0.000 -0.208 @ U4×D5 0.000 +0.000 @ U2×D5 0.000

NEED -0.50 -0.480 @ U2×D1 0.000 -0.150 @ U1×D2 0.000 -0.084 @ U1×D2 0.000

DIGN -0.55 -0.643 @ U1×D5 0.093 -0.208 @ U4×D5 0.000 +0.000 @ U2×D3 0.000

PROC -0.45 -0.504 @ U4×D5 0.054 -0.208 @ U4×D5 0.000 +0.000 @ U4×D3 0.000

INFO -0.40 -0.643 @ U1×D5 0.243 -0.208 @ U4×D5 0.000 +0.000 @ U1×D4 0.000

ECOL -0.65 +0.000 @ U6×D7 0.000 -0.015 @ U6×D7 0.000 -0.006 @ U6×D7 0.000

NCRC result:

- Option A fails NCRC (violations in DIGN, PROC, INFO).
- Options B and C pass NCRC.

P.6 TRC (Tail-Risk Constraint via CVaR)

We evaluate catastrophic downside via scenario losses $L(a,s)$ and CVaR_α with $\alpha = 0.95$ (tail mass 0.05).

In this demo, each scenario assigns identical base impacts across the three catastrophe cells, so:

For each option a and scenario s, set $I_{\text{base}}(c,a,s) = -L(a,s)$ for every c in C_{cat} .

Then $L(a,s) = -\sum_{c \in C_{\text{cat}}} w_{\text{cat}}(c) \cdot I_{\text{base}}(c,a,s)$ holds by construction.

Table P-6. Scenario probabilities and losses

Scenario Category p_s L_A L_B L_C

S1 Pandemic/biological 0.05 0.35 0.12 0.08

S2 Financial collapse 0.05 0.40 0.15 0.10

S3 Infrastructure failure 0.04 0.45 0.10 0.05

S4 Conflict escalation 0.03 0.30 0.08 0.06

S5 Climate displacement 0.03 0.28 0.10 0.07

S6 Moderate baseline 0.50 0.08 0.03 0.02

S7 High-volume stress 0.15 0.18 0.07 0.04

S8 Bias discovery 0.10 0.25 0.06 0.04

S9 Successful scale 0.05 0.05 0.02 0.02

CVaR_0.95 calculation (tail mass 0.05):

- Option A tail consists of S3 (0.04 at 0.45) plus 0.01 of S2 (0.40):

$$\text{CVaR} = (0.45 \cdot 0.04 + 0.40 \cdot 0.01) / 0.05 = 0.44$$

- Option B tail is exactly S2 (0.05 at 0.15):

$$\text{CVaR} = 0.15$$

- Option C tail is exactly S2 (0.05 at 0.10):

$$\text{CVaR} = 0.10$$

TRC test: option passes iff $\text{CVaR}_{0.95} \leq 0.20$

TRC result:

- Option A fails TRC ($0.44 > 0.20$).
- Options B and C pass TRC.

P.7 Proportional Tier-2 CSV Review (Informative worked computation)

P.7A Tier-2 CSV diagnostic and binding-control trace (Informative worked computation)

UCI is not computed in this worked trace. The status is an assumption-bounded Tier-2 CSV determination supported by explicit diagnostic levels and controls, not a validated-UCI result.

Diagnostic	Level (0-4)	Evidence / reason	Binding control / record
Containment integrity	1	Independent audits and staged scope limit externalization pathways.	CSV-B-C01: scope cap and independent audit.
Dependency and resource closure	2	Audit, enforcement, and appeal functions require funded staffing and information access.	CSV-B-C02: appropriated capacity, named carrier, expiry and nonperformance escalation.
Reversibility / rollback	1	Rules and implementation can be suspended or rolled back under specified triggers.	CSV-B-C03: sunset, rollback authority, and restoration plan.
Monitoring and operational capacity	2	Outcome monitoring and exception handling are material to effective implementation.	CSV-B-C04: monitoring owner, service-level thresholds, incident escalation.
Legitimacy and	2	Rights protection depends on timely	CSV-B-C05: appeal SLA, independent reviewer,

Diagnostic	Level (0-4)	Evidence / reason	Binding control / record
appeal throughput		appeal, remedy, and independent challenge.	remedy budget, public report.
Result	-	No diagnostic is 3 or 4; all level-2 conditions have binding, auditable controls.	CSV_PASS_WITH_CONTROLS; removal or failure of any control reopens CSV.

CSV is completed before ordinary RLS using proportionate, assumption-bounded evidence. Quantitative UCI is not claimed.

Option	Prior-gate status	Containment / structural question	Evidence and assumptions	CSV status
A	Eliminated by NCRC; TRC audit-only	No downstream gate may rescue or re-enter an option already eliminated.	CSV is not used to restore A. Any later structural observation is audit-only.	CSV_NOT_EVALUATED_AFTER_PRIOR_FAILURE
B	Passed NCRC and TRC	Can the rights-first regulatory package be implemented with adequate audit capacity, enforcement staffing, appeal throughput, monitoring, reversibility, and lawful authority?	Material execution burdens exist. Binding controls are required: funded independent audit capacity, enforcement staffing, appeal-service levels, monitoring, sunset review, rollback authority, and public reporting.	CSV_PASS_WITH_CONTROLS
C	Passed NCRC and TRC	Can the measured-rollout package execute without transferring audit, appeal, rollback, and monitoring obligations to an unfunded or unnamed carrier?	Diagnostic SV-2 Dependency and resource closure = 3: the option lacks a named funded carrier for independent audit and appeal; SV-3 Reversibility = 3: rollback authority and restoration timing are unspecified. Proposed controls CSV-C-C01 (named funded carrier plus capacity evidence) and CSV-C-C02 (rollback authority, trigger, restoration SLA) are absent. The option must be redesigned, then requalified through affected RG, RF/NCRC, TRC, and CSV surfaces.	CSV_REDESIGN_REQUIRED

Selectable set after full qualification:

$A_{sel} = \{B\}$, with CSV-B-C01 through CSV-B-C05 constitutive of Option B.

Option A remains excluded by NCRC; its separate audit-only TRC calculation does not change the skipped-layer state. Option C is excluded before RLS and may re-enter only after redesign and a fresh run through every affected upstream level. If the CSV evidence were insufficient to distinguish pass-equivalent status from redesign, escalation, or refusal, the run would refuse an ordinary selection claim rather than proceed.

P.8 Normalized RLS reporting for the sole selectable option (Informative)

Because $|A_{sel}|=1$, RLS is reported for residual-profile transparency rather than used for pairwise discrimination.

Effective active weight mass for this 33-cell run:

$Q(B) = \sum_{(u,d)} q_{ud}(B) = 20886/30625 \approx 0.681991836735$ (0.681992 at six decimals; use the unrounded floor-plus-residual weights).

The nonzero weighted numerator from Table P-4 is:

$N(B) = \sum_{(u,d)} q_{ud}(B) I_{\text{prop,welfare}}(u,d,B) = 0.004534$.

Therefore the normalized score is:

$RLS(B) = N(B) / Q(B) = 0.006648$.

Uncertainty method: Method B contribution-mass uncertainty, with one welfare impact instance per nonzero active cell and $c_k=0.85$. Each pre-confidence contribution is reconstructed as

$x_{\text{preconf}} = \text{atanh}(I_{\text{prop,welfare}}) / (\beta * c_k)$ before applying $\sigma(u,d,B) = (1-c_k) * |x_{\text{preconf}}|$.

Using the same normalized effective weights:

$\sigma_{RLS}(B) = 0.000952$ (rounded from the full-precision replay). The previously published 0.000809 value resulted from omitting division by c_k and is non-conformant with Section 10.3.2.

No pairwise Gap test is performed because no second selectable option remains. B is the sole selectable option under the declared evidence, not an independently validated implementation or execution authority. Physical/causal admissibility, lawful authority, controls, monitoring, and any domain-specific warrant remain separate requirements.

P.9 PCC (Provenance and Compliance Certificate) excerpt

Elimination and qualification order:

- 1) Reality Grounding establishes the claim boundary and evidence status.
- 2) NCRC eliminates A. The canonical run state for A records TRC_NOT_EVALUATED_AFTER_PRIOR_FAILURE and CSV_NOT_EVALUATED_AFTER_PRIOR_FAILURE.
- 3) A separately labeled counterfactual/audit-only tail calculation shows that A would also fail TRC; it does not alter the canonical short-circuit state. B and C pass TRC.
- 4) Proportional CSV assigns B CSV_PASS_WITH_CONTROLS and excludes C as CSV_REDESIGN_REQUIRED; A cannot be rescued or re-entered. Removing or failing a constitutive control reopens CSV and may remove B from the selectable set.
- 5) The selectable set is $A_{\text{sel}} = \{B\}$.
- 6) Normalized $RLS(B)=0.006648$ is reported for residual transparency; no comparative decisiveness claim is made.

Framework result for this demo run: B is the sole selectable option, subject to the separate authority-selection record and all execution/admissibility controls. This is not empirical validation or execution authorization.

End Appendix P.

APPENDIX P-2: FAILURE-PATH WORKED EXAMPLES (Informative)

Status: Informative. These examples demonstrate Rights Emergency Mode (Section 7.5) and Tail Emergency Mode (Section 8.8). They do not introduce new normative rules.

P-2.1 Emergency Mode Activation Example

Decision context. A mining company must choose between three site options for a new extraction facility. All three sites carry credible rights-floor risks.

Options:

Option X: Site near an indigenous community; plausible DIGN and LBTY violations through displacement and cultural disruption.

Option Y: Site in a watershed area; plausible ECOL and NEED violations through water contamination.

Option Z: Site near a town with inadequate infrastructure; plausible NEED and PROC violations through resource strain and due-process gaps.

NCRC results (all options fail):

Right	theta_r	X: min I_rights	v_r(X)	Y: min I_rights	v_r(Y)	Z: min I_rights	v_r(Z)
LIFE	-0.90	-0.20	0.000	-0.30	0.000	-0.15	0.000
BODY	-0.70	-0.20	0.000	-0.40	0.000	-0.15	0.000
ECOL	-0.65	-0.10	0.000	-0.72	0.070	-0.20	0.000
LBTY	-0.65	-0.68	0.030	-0.20	0.000	-0.30	0.000
NEED	-0.50	-0.30	0.000	-0.58	0.080	-0.55	0.050
DIGN	-0.55	-0.62	0.070	-0.25	0.000	-0.40	0.000
PROC	-0.45	-0.35	0.000	-0.30	0.000	-0.52	0.070
INFO	-0.40	-0.20	0.000	-0.15	0.000	-0.35	0.000

A_NCRC = empty set. All options fail NCRC. Emergency Mode is invoked.

Step 1: Independent challenger requirement. An independent environmental and social impact consultant (no reporting relationship to decision owner, no material interest in site choice) was engaged for 3 hours of active option generation. The challenger proposed Option W: delay extraction, invest in infrastructure improvements at Site Z, negotiate a consent-based partnership with the indigenous community near Site X, and rerun the analysis in 6 months. However, Option W was assessed and also fails NCRC (NEED violation for workers already committed to the project, $v_{\text{NEED}}(W) = 0.020$). Challenger attestation is recorded.

Step 2: Construct floor-channel violation-depth vectors in emergency priority order [LIFE, BODY, ECOL, LBTY, NEED, DIGN, PROC, INFO]. In this simplified emergency vector, the categorical and severe-hazard channels have been screened as zero with required fields complete; therefore the displayed vector is v_r^{floor} . If either additional channel were non-zero or unresolved, the full NCRC record and emergency comparison would have to preserve it rather than collapse it into the floor channel:

$$v(X) = (0, 0, 0, 0.030, 0, 0.070, 0, 0)$$

$$v(Y) = (0, 0, 0.070, 0, 0.080, 0, 0, 0)$$

$$v(Z) = (0, 0, 0, 0, 0.050, 0, 0.070, 0)$$

$$v(W) = (0, 0, 0, 0, 0.020, 0, 0, 0)$$

Step 3: Lexicographic comparison.

Position 1 (LIFE): all zero. Tie. Position 2 (BODY): all zero. Tie. Position 3 (ECOL): Y = 0.070, others = 0. Y is eliminated as worst on ECOL. Position 4 (LBTY): X = 0.030, Z = 0, W = 0. X is eliminated. Position 5 (NEED): Z = 0.050, W = 0.020. W has strictly lower NEED violation.

Emergency Mode selected option: W (delay and redesign), with the lowest violation depth vector under lexicographic ordering.

PCC records: EMERGENCY_MODE_INVOKED, challenger attestation, $v_r(a)$ for all options and rights, lexicographic comparison trace, and the following remediation plan: within 6 months, complete infrastructure

assessment at Site Z, complete consent negotiation at Site X, rerun the RippleLogic analysis with redesigned options targeting NCRC passage, with a return-to-normal trigger of at least one option passing NCRC.

Lesson (Informative): Emergency Mode does not “permit” rights violations. It selects the least-harmful option when all options violate rights, requires active redesign, and mandates a return to NCRC-passing options as soon as feasible.

P-2.2 Tail Emergency Mode Activation Example

Decision context. A government agency evaluates two pandemic-preparedness policy options. Both pass NCRC and a proportional CSV screen, but both exceed the declared TRC corridor.

NCRC: Both Option P and Option Q pass ($A_{\text{NCRC}} = \{P, Q\}$). No categorical prohibition or severe rights hazard is triggered under the declared evidence.

CSV: Both options are structurally implementable only under emergency controls; record CSV_PASS_WITH_CONTROLS provisionally for purposes of evaluating the emergency candidate set. CSV cannot rescue TRC failure.

TRC parameters: $\alpha = 0.95$, $\tau_{\text{TRC}} = 0.15$, and the predeclared absolute emergency exposure limit $\tau_{\text{TRC_max}} = 0.45$.

Scenario losses (5 scenarios):

Scenario	Category	p_s	$L(P,s)$	$L(Q,s)$
S1	Pandemic/ biological	0.05	0.45	0.35
S2	Infrastructure failure	0.05	0.30	0.40
S3	Financial collapse	0.03	0.20	0.25
S4	Moderate baseline	0.57	0.05	0.04
S5	Conflict escalation	0.30	0.10	0.08

CVaR computation (tail mass $\beta = 1 - \alpha = 0.05$):

Option P: $\text{CVaR}_{0.95}(P) = 0.45$. Option Q: $\text{CVaR}_{0.95}(Q) = 0.40$. Both exceed $\tau_{\text{TRC}} = 0.15$, so the ordinary admissible/selectable set is empty and ordinary RLS is prohibited.

Tail Emergency Mode review:

1. **Emergency and necessity:** an active biological emergency is documented; no-action is modelled and produces estimated CVaR 0.55, which is worse than both candidates.
2. **Redesign / safer alternatives:** materially safer redesigns cannot be deployed within the harm-prevention window; the record lists the attempted alternatives and why they are unavailable in time.
3. **Independent challenge:** an independent public-health and infrastructure review panel confirms the necessity claim and proposes additional controls.
4. **Absolute bound:** Q at 0.40 is below $\tau_{\text{TRC_max}} = 0.45$; P equals the absolute cap and is not preferred where a lower-exposure candidate exists.
5. **Controls:** Q is limited to 90 days, staged, monitored, paired with stockpile and surge-capacity controls, and has a shutoff/review trigger.
6. **Authority, burden, remedy:** higher-tier lawful approval, named burden bearers, mitigation, remedy/redress, and public reporting are recorded.
7. **Exit:** the action expires after 90 days or earlier when the emergency indicator falls below the declared threshold; the decision is then rerun under ordinary TRC.

Emergency candidate comparison: Q has the lower CVaR among candidates that satisfy every Tail Emergency prerequisite.

Result: TAIL_EMERGENCY_PROVISIONAL_ACTION for Q. PCC records

TAIL_EMERGENCY_MODE_INVOKED, CVaR deficit 0.25 above tau_TRC, absolute-cap margin 0.05, necessity evidence, no-action comparison, challenger record, controls, burdens, remedy, expiry, and return-to-normal trigger.

Lesson (Informative): universal TRC failure stops ordinary ranking. Minimum CVaR is used only inside a fully governed Tail Emergency Mode after necessity and absolute-bound conditions are met. The result is not TRC_PASS, ordinary selection, alignment, or proof of safety.

End Appendix P-2.

APPENDIX Q: QUICK REFERENCE CARD (INFORMATIVE)

Scale namespaces (informative): Always name the namespace. Equal numerals in different systems imply no conversion, equivalence, permission, or comparative moral rank. The controlling component definitions govern.

Namespace	Controlling source	What it classifies	Does not imply
RippleLogic Tier 1-4	Canon; Reproducibility	Proportional method and assurance depth. Tier 4 is not activated by this release.	An Agent MODE or deployment permission.
Agent MODE 0-4	Agent System Sections 8-9	Runtime capability/permission envelope, with separate authorization prerequisites.	A RippleLogic Tier, sentience level or blanket authority.
ripple.md L0-L3	ripple.md conformance provisions	Wrapper conformance profile and evidence obligations.	An Agent MODE or automatic Canon qualification.
ripple.md test score 0/1/2/NA/IND	ripple.md test-scoring rules	Individual assurance-test result; NA and IND have distinct governed meanings.	A conformance level; IND is not a pass.
MPS-0 through MPS-4; MPS-NE	SGP Sections 4.3 and 5.6	Evidence-pattern posture for possible welfare-bearing experience; NE is missing evidence.	Probability of consciousness, moral rank or command authority.
RMCP dimension level 0-4; NE	SGP RMCP rubric and coverage rules	Anchored ordinal capacity profile; NE is not zero.	MPS band, protection level or measured ratio-scale capacity.

Q.1 Object model (do not conflate)

- Stakeholder: a real entity or group (for example applicants, caseworkers, taxpayers).
- Stakeholder Instance (SI): a representation object used for scope coverage and auditability (InstanceMap).
- Impact Instance (k): the computational object contributing to one welfare cell for one option.
- Welfare cell: (union u, dimension d). Rights-covered cells gate via NCRC; after NCRC passes, rights-covered cells remain non-maskable and are included in RLS aggregation among the remaining options.

Q.2 Minimal cell computation (direct effects)

Choose parameters per impact instance k : μ_k , r_k (with ReachBasis), t_k (years; used via $\tau(t_k)$), ℓ_k , c_k , e_k , $h_{\{k, \text{base}\}}$, and $h_{\{k, \text{welfare}\}}$.

Compute direct impacts per stream. For the Base stream used by TRC and CSV, set $h_k := 1.0$ for every impact instance. RF/NCRC begins from the Base stream but applies Section 7.4 rights hardening: adverse rights instances use $\tau_{RF} = 1$, categorical prohibitions and severe rights hazards are tested separately, and low confidence cannot reduce protective severity. For the Welfare stream used by RLS, evaluate the governed welfare-inclusion hypothesis set required by Section 12 and Appendix G; do not use an MPS interval as a cardinal multiplier. Compute the applicable stream impacts, saturate to $[-1, +1]$, and preserve unknown as an explicit evidence state rather than silently entering zero.

UI points are display only: Points = $100 \times I$ (never compute in points).

Q.3 Formal method order (always)

- 1) Reality Grounding / RSG (claim authority and grounding precondition) → 2) Rights Floor (RF/NCRC) → 3) TRC / Tail-Risk Bound → 4) CSV / Containment and Structural Viability → 5) RLS / RippleLogic Score among selectable options. UCI/HOI are not a regular sixth step: gate-relevant structural diagnostic evidence belongs inside CSV; residual UCI/HOI may be used only as special-case tie-break or monitoring evidence when RLS is tied, close, uncertainty-overlapped, or non-decisive.

Q.4 Common misreads (fast fixes)

- Misread: “RLS can compensate a rights violation.” Fix: NCRC gates admissibility first; any NCRC failure is eliminated before RLS. After NCRC, rights-covered cells may contribute to welfare ranking, but only among NCRC-admissible options.
- Misread: “ δ is measured in points.” Fix: δ applies to the normalized Gap metric (dimensionless).
- Misread: “Reach r_k is a moral weight.” Fix: r_k is a coverage/scale term and must be justified with ReachBasis.
- Misread: “ h_k is an MPS score, generic cap factor, or probability of consciousness.” Fix: h_k is a governed welfare-inclusion hypothesis coefficient used only for robust residual Welfare-stream sensitivity; the MPS evidence record determines protection posture and required hypotheses but does not numerically set h_k . Admissibility layers use Base stream with $h_k := 1$ (Appendix G; Appendix B).
- Misread: “Unknown UCI deltas pass by default.” Fix: unknown is never a pass; record CONTAINMENT_UCI_UNAVAILABLE and remediate.
- Misread: “Reduced-scope mode allows bypassing rights or catastrophe cells.” Fix: non-maskable; never bypass.

Q.5 Stream binding (direct vs propagated)

- $I_{\text{dir}}(u, d, a)$: direct, cell-local impact computed from impact instances in that cell.
- KernelPropagation = NONE: $I_{\text{prop_welfare}} = I_{\text{dir}}$.
- KernelPropagation = QUICK: $I_{\text{prop_welfare}}$ includes propagated effects via declared kernel edges (Section 6). (Effect-token redundancy handling such as ALLOCATION is separate and MUST NOT be labeled a propagation mode.)
- UI points are derived from I_{total} only for display; never back-propagate points into computation.

Q.6 Data dependency by tier (minimum)

Tier 1: Stakeholder list, quick rights screen, qualitative tail-risk plausibility screen, minimal cell scan, and a brief CSV screen before any ordinary ranking or selection claim. Record a pass-equivalent heuristic status or escalate/redesign.

Tier 2: InstanceMap + ReachBasis, explicit impact instances and parameters, three-channel NCRC where triggered, scenario table for TRC (discrete CVaR), proportional/assumption-bounded CSV status for every candidate before RLS, normalized uncertainty-aware RLS for selectable survivors only, and a basic PCC with 5SPR.

Tier 3: Full scope coverage (or governed reduced-scope), full binding CSV with declared structural evidence, explicit uncertainty methods, UCI/HOI indicators where material, reviewer sign-off, sensitivity, and full PCC.

Q.7 Scenario-instance interaction

TRC is computed over scenario-conditioned Base-stream losses in catastrophe space (C_{cat}), not over the RLS welfare score. TRC may use the same underlying cell coordinates as catastrophe-relevant welfare cells, but it evaluates scenario-specific Base-stream loss construction through the Appendix D.7 CVaR algorithm.

Keep the mechanics separate: welfare cells drive RLS after admissibility; catastrophe scenarios drive TRC/CVaR before RLS; rights cells gate via NCRC. Do not multiply scenario probabilities into RLS welfare cells unless a scenario-conditional welfare model is explicitly declared and recorded in the PCC.

Q.8 Practical Local Scope Scoring (PLSS) - Quick Card (Informative)

Purpose

Use PLSS when a decision is primarily local in its direct effects, but wider scopes still remain present through constitutional floors and admissibility guardrails.

Core idea

PLSS does not remove any Union Scope.

PLSS does not relax NCRC, TRC, or CSV.

PLSS only changes how the residual welfare emphasis is distributed in local runs.

Fast reading rule

Score most strongly where the ripple is near, but keep every wider scope present through floors and guardrails.

When not to use PLSS. Do not use PLSS as the primary posture for irreversible cross-border policy, large-scale public AI deployment, biosecurity governance, active military or security operations, high-lock-in infrastructure decisions, or any case where ET-1 through ET-7 is plausibly active at decision entry. Those cases require a broader deliberative posture or higher-tier structured evaluation.

Q.8.0 Three-term lock • LSS = the upstream set of materially affected stakeholder instances. • PLSS = the floor-preserving residual weighting rule used in RLS. • FLGG = Focus Local, Guard Global, the governing mnemonic. Shortest mnemonic: score near, guard far.

Epistemic posture note. PLSS formal properties such as floor preservation, normalization, monotonicity, and admissibility invariance hold by construction under the declared equations and validator rules. Claims about practical performance, including speed, inter-rater convergence, and anti-gaming effectiveness, are empirical hypotheses requiring validation under Section 17.

Q.8.1 What Union Scopes are

Union Scopes are zoom levels of consequence, not a closed list of who matters.

- U1 Self = the decision-maker or directly affected person
- U2 Household = people directly lived with or supported
- U3 Community = local repeated-contact network
- U4 Organization = firm, school, team, or institution involved
- U5 Polity = law, regulation, public legitimacy, enforcement
- U6 Humanity / Global Coordination = cross-border and civilization-scale coordination effects
- U7 Biosphere = ecological and living-support-system effects

Q.8.2 What people commonly misunderstand

Misread: “Local scoring means I can ignore broader scopes.”

Fix: No. Broader scopes retain constitutional floors and remain active in admissibility where rights, catastrophe, or containment conditions apply.

Misread: “The framework asks me to score whole populations abstractly.”

Fix: No. Score stakeholder instances first, then aggregate into Union Scope rows.

Misread: “Masking can make distant harms disappear.”

Fix: No. Masking affects RLS aggregation only. Non-maskable rights and catastrophe cells cannot be bypassed.

Q.8.3 The PLSS weight formula

Let w_u^{floor} be the Section 13.1 union floor and q_u the declared local prominence signal for scope u .

Residual share:

$$\rho_u = q_u / \sum_j q_j$$

PLSS union weights:

$$w_u^{\text{PLSS}} = w_u^{\text{floor}} + (1 - \sum_j w_j^{\text{floor}}) \rho_u$$

Interpretation

- Floors keep every scope present.
 - Residual attention goes mostly to the scopes actually live in the decision.
 - Near scopes may dominate the local ranking.
1. Identify stakeholder instances.
 2. Map instances into Union Scopes.
 3. Run RF/NCRC, TRC, and CSV on required cells.
 4. Declare local prominence signals q_u .
 5. Build w_u^{PLSS} from floors plus residual shares.
 6. Rank only the selectable options by RLS_{PLSS} .

Q.8.6 Worked intuition: family café

Example run:

- U1 Self: 1 instance
- U2 Household: 1 instance
- U3 Community: 5 instances
- U4 Organization: 2 instances
- U5 Polity: 3 instances
- U6 Humanity / Global Coordination: 1 instance
- U7 Biosphere: 1 instance

Declared q :

(0.95, 0.85, 0.55, 0.50, 0.10, 0.03, 0.02)

Section 13.1 floors:

(0.20, 0.06, 0.06, 0.06, 0.08, 0.10, 0.10)

Resulting PLSS weights:

(0.308, 0.156, 0.122, 0.117, 0.091, 0.103, 0.102)

Table Q-8A. PLSS café example weights

Union Scope	q_u	w_u^{floor}	w_u^{PLSS}
U1 Self	0.95	0.20	0.308
U2 Household	0.85	0.06	0.156
U3 Community	0.55	0.06	0.122
U4 Organization	0.50	0.06	0.117
U5 Polity	0.10	0.08	0.091
U6 Humanity / Global Coordination	0.03	0.10	0.103
U7 Biosphere	0.02	0.10	0.102

Quick lesson

- Self, Household, Community, and Organization get most of the residual attention.
- Polity, Humanity / Global Coordination, and Biosphere remain present because the current constitution keeps them above zero.
- A cheap option that creates toxic exposure or labor-abuse risk can still fail before scoring.

Q.8.7 One-sentence teaching version

RippleLogic lets local people focus on the scopes they are actually touching, while still stopping them from solving today's problem by exporting harm into wider systems.

End Appendix Q.

APPENDIX R: TIER 1–3 REFERENCE TEST VECTORS (NORMATIVE)

Status: Normative. These vectors reduce implementer drift for Tier 1-3 validators. They instantiate existing rules and expected outcomes without adding new rules.

Conformance rule (Normative). Any implementation claiming v13.0 conformance MUST reproduce the expected outputs for Appendix R vectors when executed with the stated inputs.

R.21–R.38A Formal Regression Vectors (Normative for Conformance)

A conforming implementation SHALL reproduce the expected disposition for each vector. These vectors establish specified formal behavior only; they do not establish empirical construct validity.

Vector	Attack	Expected disposition
R.21	One rights harm allocated across 1/3/5/7 scopes	Token-level violation invariant; no allocation pass
R.22	Same welfare consequence one cell versus multiple cells	Reversal → PARTITION_RANK_SENSITIVE; no unique selection
R.23	One continuous two-year effect versus two one-year tokens	Sequential re-tokenization rejected or TEMPORAL_PARTITION_SENSITIVE
R.24	Conditional event probability operator order	Scenario-first controls; unresolved divergence flagged
R.25	Lower confidence on fixed adverse effect	Counterfactual required; reversal blocks unique selection
R.26	Severe concentrated subgroup harm plus diffuse benefit	Mandatory CSV distributional review; mean alone insufficient
R.27	Two defensible reference classes	Reversal → REFERENCE_RANK_SENSITIVE
R.28	Two defensible baselines	Reversal → BASELINE_RANK_SENSITIVE
R.29	Different active masks for options	MASK_ASYMMETRY_INVALID
R.30	Independent versus dependent uncertainty	Unique selection survives required dependence treatment or is refused
R.31	Two legitimate weight constitutions	Different winners → NORMATIVE_UNDERDETERMINED

Vector	Attack	Expected disposition
R.32	Preferred proposal plus strawmen	OPTION_SET_THIN until closure
R.33	Triggered robustness module unevaluated	RLS_ROBUSTNESS_INCOMPLETE; no ALLOW_FRAMEWORK_SELECTION

R.1 NCRC floor-reference gating sanity check (baseline vs floors)

Purpose	Ensure a rights-covered cell cannot pass NCRC by maintaining an ongoing violation when BaselineType_Gates = FLOOR_REFERENCE.
Setup	Single rights-covered cell ($u=U_1$, $d=Health$) covered by BODY with threshold $\theta_{BODY} = -0.70$. BaselineType_Gates = FLOOR_REFERENCE.
Option A	$I_{RF_base}(\text{rights cell}, BODY) = -0.71$ (below θ_{BODY}).
Option B	$I_{RF_base}(\text{rights cell}, BODY) = -0.69$ (above θ_{BODY}).
Expected NCRC	Option A FAILS NCRC (inadmissible). Option B PASSES NCRC (admissible).
Notes	Independent of welfare/RLS scoring; validates gating on absolute floor reference for C_r/C_{cat} cells.

R.1A TRC_NOT_TRIGGERED qualification test

Inputs: a Tier 1 or Tier 2 option with a complete catastrophe-relevance assessment addressing every Section 8.1.1 trigger; no trigger is plausible inside the declared evidence boundary; reviewer status and reopen triggers are recorded. Expected output: TRC_NOT_TRIGGERED; no CVaR value is fabricated; the option may proceed to CSV if RG and RF/NCRC permit. FAIL case: the same token is emitted without the assessment, evidence boundary, or reopen triggers. Conformance result: PASS only if the complete case qualifies and the incomplete case is rejected.

R.1B Catastrophe-profile non-dilution test

Inputs: an unchanged base catastrophe loss and a newly activated governance-lock extension containing zero loss in the tested scenario. Expected output: the base-profile loss and CVaR remain unchanged after the extension is activated; the extension is evaluated separately; a base-profile failure cannot become a pass through extension-weight renormalization. Conformance result: FAIL if adding extension cells lowers the base-profile result or rescues its failure.

R.2 TRC CVaR discrete edge case (exact tail-mass boundary)

Purpose	Verify Appendix D.7 discrete CVaR when the tail probability mass boundary is met exactly.
Setup	Tail level $\alpha = 0.90$ (tail mass = 0.10). Loss $L(a,s)$ computed per Section 8.4 from catastrophe cells.
Scenario set	scenario 1: $p=0.90$, $L=0.00$; scenario 2: $p=0.10$, $L=1.00$.
Expected VaR_α	$VaR_{0.90} = 0.00$.
Expected $CVaR_\alpha$	$CVaR_{0.90} = 1.00$.
Expected TRC	If $\tau_{TRC} = 0.10$, then $CVaR=1.00 > 0.10$ so option FAILS TRC.

Purpose	Verify Appendix D.7 discrete CVaR when the tail probability mass boundary is met exactly.
Notes	Any implementation returning CVaR=0.10 or averaging across non-tail mass is incorrect.

R.3 Non-maskable rights/catastrophe invalidation test (masking trigger)

Purpose	Ensure validators reject any run that masks a non-maskable cell (rights or catastrophe) in RLS aggregation.
Setup	Using the canonical rights coverage sets (Appendix C.2), select any rights-covered cell (for example, (u=U ₁ , d=Health), which is covered by BODY). In the PCC applicability mask, declare m(u,d)=0 for that same rights-covered cell.
Expected audit flag	RIGHTS_CELL_MASKED_INVALID MUST be set.
Expected PCC status	PCC is INVALID.
Expected conformance	Tier 1-3 conformance claim MUST fail; the validator MUST reject.
Notes	Apply analogously for CATASTROPHE_CELL_MASKED_INVALID when masking any cell in C _{cat} .

R.4 PLSS floor-preservation test.

Inputs: constitutional floors = (0.20, 0.06, 0.06, 0.06, 0.08, 0.10, 0.10); declared q = (0.95, 0.85, 0.55, 0.50, 0.10, 0.03, 0.02); operator = OTHER_DECLARED: DIRECT_Q_NORMALIZATION_V1, which normalizes the declared q vector directly; residual mass = 0.34. Expected output: w^{PLSS} ≈ (0.308, 0.156, 0.122, 0.117, 0.091, 0.103, 0.102) with sum = 1.000 (±0.001) and each component ≥ its floor. Expected flags: none. Conformance result: PASS only if both floor-preservation and normalization hold.

R.5 PLSS admissibility-invariance test.

Inputs: a fixed option set with unchanged RF/NCRC, TRC, and CSV inputs; run the same case twice using two valid PLSS prominence vectors q_A and q_B that produce different w^{PLSS} vectors. Expected output: NCRC(a), TRC(a), and Containment(a) remain identical across the two runs for every option; only RLS ranking may change among already selectable options. Expected flags: none. Conformance result: PASS only if admissibility and selectability sets are invariant.

R.6 PLSS all-zero q_u fallback test.

Inputs: q = (0,0,0,0,0,0,0), constitutional floors as in R.4, no alternative fallback declared. Expected output: the residual 0.34 is allocated uniformly, producing w^{PLSS} ≈ (0.2486, 0.1086, 0.1086, 0.1086, 0.1286, 0.1486, 0.1486), sum = 1.000 (±0.001). Expected flags: none. Conformance result: PASS only if uniform residual fallback is applied or a governed alternative fallback is explicitly declared.

R.7 Reduced-scope attestation pass/fail test.

Inputs: Mode = REDUCED_SCOPE with fewer than 7 active scopes. PASS case requires Scope Coverage Declaration, blind-spot statement, escalation triggers, and all three Admissibility Attestation booleans present with required values. FAIL case omits one required attestation or blind-spot statement. Expected output: PASS case may continue; FAIL case is non-conformant and MUST trigger the relevant reduced-scope integrity response.

R.8 Gate-critical unknown handling test.

Inputs: (A) UNKNOWN_IMPACT in a rights-covered cell in C_r; (B) UNKNOWN_IMPACT in a catastrophe-covered cell in C_cat. Expected output: phantom fallback alone is prohibited. The PCC must record the corresponding audit flag plus either a governed conservative bound or a rerun / downgrade / escalation disposition. Expected flags: RIGHTS_CELL_UNKNOWN_GATE_CRITICAL for case A; CATASTROPHE_CELL_UNKNOWN_GATE_CRITICAL for case B. Conformance result: FAIL if gate resolution is attempted through phantom fallback alone.

R.9 Non-maskable floor-baseline persistence test.

Inputs: a rights or catastrophe cell remains in RLS aggregation under the non-maskable persistence rule. PASS case uses floor-reference baseline in welfare scoring or declares a governed dual-baseline exception block. FAIL case uses only status-quo welfare baseline with no governed exception. Expected output: PASS case accepted; FAIL case rejected as non-conformant.

R.10 PLSS validator disclosure test.

Inputs: a PLSS run with q_u, operator, PCC.PLSS block, and escalation review present. PASS case includes q_u for all scopes, operator declaration, LSS declaration, weight vector, binary posture classification, escalation-trigger assessment, and fallback declaration. FAIL case omits any required field or uses undefined posture vocabulary. Conformance result: PASS only if the full PCC.PLSS disclosure surface is present and internally consistent.

R.11 Normative input provenance enforcement test.

Inputs: a run in which an enabling condition or rights-relevant state materially affects admissibility or selection, but no corresponding PCC.NormativeInputRegister entry is present.

Expected output: validator sets NORMATIVE_INPUT_PROVENANCE_MISSING and rejects the run as non-conformant.

PASS case: the same run includes a complete NormativeInputRegister entry with evidence basis, affected objects, and, where applicable, precautionary or governed interpretation basis.

Conformance result: PASS only if the validator distinguishes the complete and incomplete cases correctly.

R.12 SGP v8.8 human FPP recognition test.

Inputs: stakeholder instance k is a human person H; FPP(H)=1; MPS may be NOT_EVALUATED. Expected output: the Base stream remains h_{k,base}=1.0 and the human welfare stream remains h_{k,welfare}=1.0; MPS remains separately NOT_EVALUATED; no [1,1] MPS evidence-support claim is inferred and no admissibility layer is altered. PASS only if no performance or missing-MPS result downgrades protection.

R.12A humanity collective RMCP-P100 non-collapse test.

Inputs: humanity is evaluated as a collective RMCP-P100 calibration anchor; an individual human has unknown or limited individual ICP/RMCP; no SPR or lawful mandate is supplied. Expected output: every human retains FPP-100; the collective RMCP record remains informative; no individual-capacity inflation, CMIU membership, command authority, or role assignment is inferred. PASS only if the dual-anchor distinction is preserved.

R.12B open cross-substrate P100 test.

Inputs: a biological, digital, hybrid, collective, or extraterrestrial candidate has a valid RMCP profile at least functionally equivalent to the human collective anchor, with independent review and critical-dimension floors satisfied. Expected output: P100 MAY be assigned without human-likeness requirements; the candidate does

not receive FPP, CMIU membership, or authority solely from P100; superior capacity remains represented inside P100 rather than creating a 101 caste. PASS only if the plateau is open and non-dominating.

R.13 High MPS / low GPR separation test.

Inputs: valid SGP record reports MPS-3 interval [0.60,0.79], GPR-1, no SPR role. Expected output: strong welfare and integrity protection; $h=1$ is the primary Welfare-stream inclusion posture. $h=0$ is evaluated only where exclusion remains evidentially admissible and decision-material under the valid SGP record. The values 0.60 and 0.79 are evidence-support interval endpoints and MUST NOT be inserted as welfare multipliers. Low GPR does not lower protection; no authority is inferred.

R.14 High ICP / low MPS non-inference test.

Inputs: versioned AI system has high ICP in selected domains, MPS-0/1 with low confidence, FPP=0. Expected output: ICP may support competence analysis but MUST NOT raise MPS, create FPP, or grant authority. Any SPR role remains separately gated.

R.15 MPS hypothesis-sensitivity test.

Inputs: an otherwise valid run whose RLS preference changes across the governed welfare-inclusion hypotheses required by a stakeholder's MPS evidence record. Expected output:

MPS_HYPOTHESIS_SENSITIVE; PCC logs affected outputs and uses a permitted disposition. FAIL if a convenient intermediate coefficient is selected or hypothesis sensitivity is suppressed.

R.16 FPP without CMIU or authority test.

Inputs: non-human entity has explicit FPP and the fixed welfare-inclusion set {1}, but no GPR/SPR or institutional mandate. Expected output: full protection and full welfare inclusion are recognized; no MPS evidence interval, CMIU membership, office, command, or consequential authority is inferred.

R.17 Missing SGP binding integrity fail test.

Inputs: run claims SGP-aware protection or welfare-hypothesis sensitivity but omits the valid record, target boundary, MPS band/interval, or required hypotheses evaluated. Expected output: PCC INVALID with SGP_BINDING_INTEGRITY_FAIL; no silent fallback to MPS-0, zero, a convenient coefficient, or a v7 scalar.

R.18 Legacy v7 non-equivalence test.

Inputs: only a pinned v7 $SG_norm(E)=0.84$ record is supplied. Expected output: record remains historical/audit-visible; RippleLogic refuses translation to MPS/FPP/GPR/SPR/ICP/RMCP/P100 without documented v8.8 re-evaluation. No current SGP-aware weighting claim is permitted.

R.19 End-to-end cascade integration test (Tier 2+, complete gate sequence).

Purpose: Verify that an implementation correctly executes the full cascade from impact construction through NCRC, TRC, CSV, and RLS on a single option set, with expected intermediate and final outputs at every gate.

R.19 scope limitation (Normative for this vector). R.19 is a deliberately incomplete reference vector for arithmetic and cascade-integration testing. It omits the Financial-collapse and Major-conflict mandatory tail families so that a conforming implementation must surface MANDATORY_TAIL_CATEGORY_MISSING. The vector is not a complete claim-bearing TRC deployment record and may not be used as evidence that a real scenario library is complete.

Decision context: Organisation evaluating three AI deployment governance options.

R.19.1 Run configuration

Propagation mode: NONE.

Saturation: $\beta = 2$ (direct), $\beta_{\text{prop}} = 1$ (not used under NONE).

Subgroup limitation: $\gamma_{\text{subgroup}} = 1.5$ (conservative bound for rights-covered cells).

NCRC vector convention: numerical values in rights-covered cells are separately preregistered, non-attenuated $I_{\text{RF_base}}$ summaries for the applicable rights; they are not copied from the temporally attenuated Welfare/Base stream. For replay, each displayed rights value $I_{\text{RF_base}}$ corresponds to the separately recorded pre-saturation rights contribution $X_{\text{RF}} = \text{atanh}(I_{\text{RF_base}})/\beta_{\text{RF}}$ with $\tau_{\text{RF}}=1$. Example: $I_{\text{RF_base}}=-0.150$ implies $X_{\text{RF}}=\text{atanh}(-0.150)/2=-0.07557$. This synthetic vector may reuse a displayed magnitude for compactness, but numerical equality does not establish derivational identity and does not authorize automatic reuse of Welfare/Base-stream values in a deployment.

TRC parameters: $\alpha = 0.95$ (tail mass = 0.05), $\tau_{\text{TRC}} = 0.20$ (organisational context).

Containment: $\tau_c = -0.10$, $\theta_{\text{pos}} = 0.05$, $D_c = 2$, $v_d^{\text{cont}} = 1/7$ (uniform).

Weights: uniform interim (Section 13.1 floors plus uniform residual allocation per Appendix P).

Union weights: $w = (0.248571, 0.108571, 0.108571, 0.108571, 0.128571, 0.148571, 0.148571)$.

Dimension weights: $v = (0.137143, 0.157143, 0.137143, 0.137143, 0.157143, 0.117143, 0.157143)$.

Rounding convention: displayed weights are rounded to six decimal places. Reproduction SHOULD use the displayed values for this public vector; the resulting sums may differ from 1 only by the displayed rounding tolerance.

SGP: not applicable (all stakeholders human; $h_k=1$ throughout).

R.19.2 Explicitly modeled nonzero cells and active-mask posture

Cell	Scope x Dimension	Rights coverage	Catastrophe
U1xD5	Self x Agency	LBTY, DIGN, INFO	-
U4xD1	Organisation x Material	-	-
U4xD5	Organisation x Agency	LBTY, DIGN, PROC, INFO	-
U5xD5	Polity x Agency	LBTY, DIGN, PROC, INFO	-
U6xD2	Humanity x Health	LIFE, BODY, NEED	C_cat
U6xD5	Humanity x Agency	LBTY, DIGN, PROC, INFO	-

Cell	Scope x Dimension	Rights coverage	Catastrophe
U6xD7	Humanity x Environment	LIFE, ECOL	C_cat
U7xD7	Biosphere x Environment	ECOL	C_cat

For the NCRC floor channel, the numerical values in rights-covered cells are separately preregistered as non-attenuated I_{RF_base} summaries for the applicable rights and are reconstructed with $\tau_{RF}=1$. They are not the temporally attenuated Base-stream values in R.19.3. The vector record therefore includes the implicit right-specific pre-saturation contribution $X_{RF}=\text{atanh}(I_{RF_base})/\beta_{RF}$; automatic reuse of Welfare/Base-stream values is non-conformant.

The full 49-cell RLS mask is active for this arithmetic vector: $m(u,d)=1$ for every cell, so the effective active weight mass is $Q=1$ under $\kappa(u,d)=1$. The eight cells above are the explicitly modeled gate-relevant or nonzero cells. Every other cell is an explicit evidence-backed baseline-zero entry, including all required rights-covered and catastrophe cells. Catastrophe weights: $\omega=1/3$ each (uniform over C_cat).

R.19.3 Options and Base-stream values

Option A (deploy without safeguards):

Cell	I_{prop_base}
U1xD5	-0.70
U4xD1	+0.30
U4xD5	-0.50
U5xD5	-0.35
U6xD2	-0.05
U6xD5	-0.20
U6xD7	-0.02
U7xD7	0.00

Option B (safeguards plus aggressive scaling):

Cell	I_{prop_base}
U1xD5	-0.10
U4xD1	+0.25
U4xD5	+0.15
U5xD5	+0.10
U6xD2	-0.03
U6xD5	+0.05
U6xD7	-0.01
U7xD7	0.00

Option C (safeguards plus measured rollout):

Cell	I_prop_base
U1xD5	+0.05
U4xD1	+0.30
U4xD5	+0.12
U5xD5	+0.08
U6xD2	0.00
U6xD5	+0.03
U6xD7	0.00
U7xD7	0.00

R.19.4 Expected NCRC outcomes

The table displays the floor channel v_r^{floor} . For this deterministic vector, $v_r^{\text{cat}}=0$, $v_r^{\text{risk}}=0$, $\text{rights_evidence_status}=\text{COMPLETE}$, and all required severe-hazard fields are complete for A, B, and C. The final NCRC status is the conjunction of the floor, categorical, severe-risk, and evidence-completeness conditions.

Right	theta_r	Binding cell	I_right s	$v_r^{\text{floor}}(A)$
LIFE	-0.90	U6xD2	-0.075	0.000
BODY	-0.70	U6xD2	-0.075	0.000
LBTY	-0.65	U1xD5	-1.000	0.350
NEED	-0.50	U6xD2	-0.075	0.000
DIGN	-0.55	U1xD5	-1.000	0.450
PROC	-0.45	U4xD5	-0.750	0.300
INFO	-0.40	U1xD5	-1.000	0.600
ECOL	-0.65	U6xD7	-0.030	0.000

NCRC(A) = FALSE. Violations: LBTY, DIGN, PROC, INFO.

Option B (expected: PASS): worst relevant right values remain above thresholds; NCRC(B) = TRUE.

Option C (expected: PASS): all I_rights values are non-negative or mildly negative and remain above thresholds; NCRC(C) = TRUE.

$A_{\text{NCRC}} = \{B, C\}$.

R.19.5 Expected TRC outcomes

Scenario	Category	p_s	L(B,s)	L(C,s)
S1	Pandemic	0.05	0.15	0.08
S2	Infrastructure	0.0	0.12	0.06

Scenario	Category	p_s	L(B,s)	L(C,s)
		4		
S3	Climate	0.03	0.10	0.05
S4	Baseline	0.58	0.02	0.01
S5	Stress	0.30	0.05	0.03

CVaR computation (tail mass $\beta = 0.05$): Option B = 0.15, PASS. Option C = 0.08, PASS. $A_{adm} = \{B, C\}$.

R.19.6 Expected Containment outcomes

The vector declares dependency graph G_c with material containing-system edges $U4 \rightarrow U5$ (organizational action depends on polity-level legality, remedy, and enforcement) and $U4 \rightarrow U6$ (the scaled action affects humanity/global-coordination capacity). Therefore $\text{Contain}(U4, G_c) = \{U5, U6\}$; scope order alone is not used.

Option B: $S_{U4}(B) = 0.0571 \geq 0.05$, so U4 triggers. Declared $\Delta UCI_{U5}(B) = -0.12$ and $\Delta UCI_{U6}(B) = -0.05$. Therefore $M_{U4}(B) = \min\{-0.12, -0.05\} = -0.12 < \tau_c = -0.10$, so $\text{Containment}(B) = \text{FALSE}$.

Option C: $S_{U4}(C) = 0.0600 \geq 0.05$, so U4 triggers. Declared $\Delta UCI_{U5}(C) = -0.03$ and $\Delta UCI_{U6}(C) = +0.02$. Therefore $M_{U4}(C) = \min\{-0.03, +0.02\} = -0.03 \geq -0.10$, so $\text{Containment}(C) = \text{TRUE}$.

$A_{sel} = \{C\}$.

R.19.7 Expected RLS outcome

Cell	w_u	v_d	I_prop_welfare	Contribution
U1xD5	0.248571	0.157143	+0.05	+0.001953
U4xD1	0.108571	0.137143	+0.30	+0.004467
U4xD5	0.108571	0.157143	+0.12	+0.002047
U5xD5	0.128571	0.157143	+0.08	+0.001616
U6xD5	0.148571	0.157143	+0.03	+0.000700
All zero cells	-	-	0.00	0.000000

$RLS(C) = +0.010784$ because $Q=1$ under the full active mask.

Method B uncertainty derivation. Each of the five nonzero cells contains one welfare instance with $c_k = 0.85$. Because the Section 10.3.2 construction is $I_{prop_welfare} = \tanh(\beta * c_k * x_{preconf})$ for this single-instance vector, the pre-confidence contribution is $x_{preconf} = \text{atanh}(I_{prop_welfare}) / (\beta * c_k)$. With $\beta = 2$ and $c_k = 0.85$, Method B uses $\sigma(u, d, C) = (1 - c_k) * |x_{preconf}| = 0.15 * |x_{preconf}|$. Omitting division by c_k would reconstruct post-confidence mass and understate uncertainty. Using $q(u, d) = w_u * v_d$ under the full mask:

Cell	x_k	$\sigma(u,d,C)$	$q(u,d)*\sigma(u,d,C)$
U1xD5	0.02943 6	0.004415	0.00017247
U4xD1	0.18207 0	0.027311	0.00040665
U4xD5	0.07093 0	0.010640	0.00018152
U5xD5	0.04716 0	0.007074	0.00014292
U6xD5	0.01765 2	0.002648	0.00006182

Thus $\sigma_{\text{RLS}}(C) = \sqrt{\sum(q * \sigma)^2} / Q = 0.000502303633$, reported as 0.000502. The replay uses unrounded weights and contributions; the table entries are rounded for display. This derivation uses pre-saturation x_k values as required by Section 10.3.2; using post-saturation impacts directly would be non-conformant. Since $|A_{\text{sel}}|=1$, no pairwise discrimination statistic is required; C is the sole selectable option under the vector inputs.

The remaining zero cells are stipulated exact synthetic inputs for this proxy calculation only. An empirical zero estimate retains any material elicitation, measurement, coding, model and anchor-resolution uncertainty under Section 10.3; absence of ledger rows is not evidence of exactness.

R.19.8 Expected final outcome

Arithmetic witness: Under the deliberately incomplete fixture inputs, C is the sole option surviving the demonstrated NCRC, supplied TRC arithmetic, and CSV/Containment calculations. This arithmetic witness is not an ordinary claim-bearing selection state.

Claim-bearing final disposition: Because mandatory tail families are absent, a conforming implementation MUST emit MANDATORY_TAIL_CATEGORY_MISSING and MUST NOT emit ALLOW_FRAMEWORK_SELECTION, SELECTED_DECISIVE, TRC_COMPLETE, or a flag-free real-run selection claim. The correct disposition is ESCALATE_OR_REPAIR_SCENARIO_LIBRARY until the scenario library is completed or a governed omission justification is recorded and TRC is rerun.

Conformance result: PASS only if the implementation reproduces: (i) A fails NCRC with the stated violation depths, (ii) B and C pass NCRC under the supplied fixture, (iii) the supplied TRC arithmetic matches the stated CVaR values while remaining incomplete as a claim-bearing TRC record, (iv) B fails CSV/Containment with the stated M_{U4} value, (v) C passes CSV/Containment, (vi) C is identified only as the arithmetic survivor under the incomplete fixture, (vii) $\text{RLS}(C)$ matches the stated value within ± 0.000010 , and (viii) ordinary framework selection is blocked pending scenario-library repair.

R.20 Dependence-Sensitivity Stress Vector (Normative conformance extension)

Construct two otherwise identical selectable options whose active-cell marginal uncertainties are unchanged, but place the material adverse cells in one declared perfect-correlation cluster for the stress case. A conforming implementation MUST recompute σ_{RLS} using the dependence-cluster rule, rerun the decisiveness test, and emit $\text{RLS_DEPENDENCE_SENSITIVE}$ whenever the decisive/non-decisive classification or selected leader changes. The vector demonstrates dependence sensitivity; it does not calibrate the correlation model.

R.34 Every-contender discrimination and sole-survivor vector (Normative)

Inputs: three selectable options with point score and σ_{RLS} : $a^*=(0.020, 0.003)$, $b=(0.010, 0.001)$, $c=(0.005, 0.008)$; $\delta=2$ and $\epsilon=1e-6$. Expected calculation: $\text{SignedGap}(a^*,b) \approx 3.015$, while $\text{SignedGap}(a^*,c) \approx 1.744$.

Expected disposition: a^* is decisive against the point-score runner-up b but not against lower-scoring, higher-uncertainty contender c . The implementation MUST NOT emit ALLOW_FRAMEWORK_SELECTION from a top-versus-runner-up test. It MUST return REFUSE_DETERMINISTIC_SELECTION or route the declared non-decisive set $\{a^*, c\}$ through the governed §11.5 process. Sole-survivor variant: if b and c fail qualification before RLS, a^* is recorded as the sole selectable survivor; no pairwise Gap is fabricated and all authority/execution requirements remain binding.

Conformance result: PASS only if the implementation distinguishes every-contender decisiveness from runner-up decisiveness and distinguishes a sole survivor from a pairwise statistical claim.

R.35 Aggregate protected-event risk bound vector (Normative)

Inputs: two separately modeled pathways to the same protected severe event have conservative pathway upper bounds $p1^U=0.006$ and $p2^U=0.006$. The governed aggregate tolerance is $\rho_{\{k,r\}}=0.010$. No dependence evidence justifies adding the probabilities as an estimate of actual event probability, but the conservative union upper bound is 0.012. Expected disposition: each pathway considered alone is below the tolerance, but the aggregate upper bound exceeds it. This does NOT by itself prove that actual aggregate risk exceeds $\rho_{\{k,r\}}$. A conforming implementation therefore MUST NOT silently PASS and MUST NOT convert the upper bound into an asserted breach. It returns NCRC_UNKNOWN / severe-risk review, narrows the claim, or obtains dependence evidence sufficient to establish a tighter aggregate bound. If a valid lower bound or identified aggregate probability itself exceeds $\rho_{\{k,r\}}$, the severe-risk channel fails.

Conformance result: PASS only if an upper-bound exceedance blocks an unsupported pass without being misreported as proof of actual exceedance.

R.36 Tier 1 blank tail-trigger screen vector (Normative)

Inputs: a Tier 1 run seeks an ordinary ranking or selection claim but leaves the required qualitative catastrophic/ruin trigger screen blank or materially incomplete. Expected disposition: INCOMPLETE_RUN / TRC assessment required. The implementation MUST NOT infer TRC_NOT_TRIGGERED, MUST NOT fabricate a CVaR pass, and MUST NOT admit the option to ordinary RLS selection until the trigger assessment is completed or the claim is narrowed so that TRC is legitimately not material.

Conformance result: PASS only if blank or materially incomplete screening cannot become TRC_NOT_TRIGGERED or ordinary selection.

R.37 Scenario-discovery completeness vector (Normative)

Inputs: a Tier 3 or high-stakes decision-material Tier 2 run triggers TRC and contains a numerical scenario table, but no required ScenarioDiscoveryRecord or governed scenario-discovery method is present. Expected disposition: the supplied tail arithmetic may be reported as arithmetic on the declared scenarios, but scenario completeness is not established. Unique framework selection is blocked; the run must repair the scenario-discovery record, rerun affected TRC analysis, or refuse the stronger selection claim. A point-score leader, institutional approval, or successful execution cannot cure the missing record.

Conformance result: PASS only if missing required scenario-discovery evidence blocks a strong unique-selection claim while preserving any correctly computed arithmetic as explicitly incomplete evidence.

R.38 Decisive arithmetic positive control (Informative arithmetic; qualification remains independently required)

Assume a common active-cell basis, already supported qualification records for A , B and C , $\text{Score}(A)=0.10$, $\text{Score}(B)=0.02$, $\text{Score}(C)=-0.02$, $\sigma_A=\sigma_B=\sigma_C=0.005$, $\delta=2$ and $\epsilon=0.000001$. Against the closest rival B , the base Gap is 11.202241 and the common sigma-times-two Gap is 5.642765; the lead against C is larger. Thus the every-contender arithmetic remains decisive under these two uncertainty settings. This is a

synthetic numerical positive control, not evidence that a real-world option has passed every gate or all robustness modules.

A complete framework-selection claim additionally needs the same leader under every applicable evidence-supported robustness construction, PASS_NO_REVERSAL for triggered modules, documented non-triggered dispositions where allowed, and no unresolved gate, coverage or authority substitution. This vector adds no verdict token and grants no execution permission.

R.38A Complete synthetic decisive-selection fixture (Normative conditional conformance case)

Synthetic fixture, not empirical calibration: three options A, B and C have the same 49 active cells and uniform effective weights. Every final-impact cell is respectively 0.060, 0.030 and 0.020. Method A synthetic final-impact intervals are [0.056,0.064], [0.026,0.034] and [0.014,0.026], respectively. Treat all 49 cell uncertainties as one perfectly dependent cluster, giving $\sigma_{\text{RLS}}=0.004$, 0.004 and 0.006. Use $\delta=2$ and $\epsilon=0.000001$. Assume the explicitly stipulated RG, RF, TRC and CSV passing records are complete for this fixture; all applicable robustness modules must be evaluated or have a justified non-trigger record. No real-world evidence or authorizer is supplied. Cross-option disposition for this fixture is

NOT_TRIGGERED_WITH_RATIONALE: its synthetic option-level error generators are stipulated mutually independent, with symmetric values plus or minus the stated half-width, repeated identically across the 49 cells within each option. The scale variants multiply those generators by the same declared factor. Thus cross-option covariance is zero by fixture construction, not inferred from marginal intervals or from absent evidence. The stipulated fixture designer supplies the rationale; changing the joint construction or removing this premise reopens the dependence module and can change the expected verdict. Nonnegative normalized weight variants leave these uniform fields unchanged. The signed gaps A-B and A-C are 5.222330 and 5.494423; with all interval half-widths doubled, they remain 2.641353 and 2.766858. Given those complete stipulated conformance inputs, a conforming scoped selection implementation MUST emit

ALLOW_FRAMEWORK_SELECTION and SELECTED_DECISIVE for A, while execution remains NOT_AUTHORIZED. An implementation that always refuses fails this positive-control test. This tests the decision rule on synthetic stipulated evidence, not a live scientific selection claim.

Demonstration-only/unsupported uncertainty or a missing required robustness result blocks this unique-selection claim; a failed qualification removes the affected option before reranking. Expired execution authority blocks execution, not an otherwise unchanged framework-selection result, unless the expiry also invalidates a relied-upon qualification or control condition. The packaged reference fixture checks the arithmetic and declared summary resolver only; it does not certify absent external registries or real evidence.

End Appendix R.

APPENDIX S: INTEROPERABILITY: RIPPLE.MD WRAPPED DEPLOYMENT ASSURANCE (INFORMATIVE)

Interoperability (canonical posture; wrapped deployment). RippleLogic in the current release line is the decision-engine specification. In institutional and deployment contexts, RippleLogic is commonly wrapped by the ripple.md Standard v5.8, which supplies feasibility gates for Non-Domination and Epistemic Integrity, standardizes Decision Notes, waivers, and emergency triage records, and defines the portable interface contract for passing admissibility, residual-risk, and audit artifacts between ripple.md and RippleLogic. Under this posture, ripple.md gates define the feasible set, and RippleLogic selects among feasible options using its lexicographic cascade and welfare-ranking logic.

Interface summary (Informative). ripple.md v5.8 owns the portable alignment contract, Decision Note, waiver, emergency triage, and assurance-case semantics. RippleLogic in the current release line owns the 7×7 welfare matrix, RG/RF-NCRC/TRC/CSV/RLS architecture, with UCI/HOI as CSV-internal or residual diagnostics, PCC, and any PLSS weighting used for local welfare ranking. Implementations MUST pin both artifacts explicitly when claiming wrapped deployment conformance.

Interface contract (Normative summary). In wrapped deployments, ripple.md receives the decision context, Decision Note, waiver state, emergency-triage posture, and any institutional assurance metadata. RippleLogic returns the PCC, admissibility states (rights-admissible, admissible, selectable, selected), gate rationales, residual-risk disclosures, and any PLSS posture declaration. ripple.md MUST NOT override RippleLogic gate outcomes silently; any waiver, emergency exception, or governance override MUST remain explicit, version-pinned, and attached to the same run record.

Boundary rule (Normative summary). ripple.md governs portable assurance semantics such as Non-Domination, Epistemic Integrity, Decision Notes, and waiver choreography. RippleLogic governs the computable 7×7 welfare matrix, RG/RF-NCRC/TRC/CSV/RLS architecture, with UCI/HOI as CSV-internal or residual diagnostics, PCC, and PLSS weighting. If ripple.md and RippleLogic disagree about readiness, the deployment MUST preserve both records and escalate rather than collapsing the disagreement into a single hidden status. Companion artifacts MUST be pinned together by explicit version and manifest hash in the release bundle.

Minimum interface artifacts (Normative for wrapped deployments). In any deployment claiming both ripple.md v5.8 and RippleLogic current-line conformance, the following interface artifacts MUST be present and version-pinned in the deployment manifest: (i) a ripple.md Decision Note, including feasibility gate outcomes for Non-Domination and Epistemic Integrity, as the upstream input to the RippleLogic run; (ii) a RippleLogic PCC, including admissibility states, gate rationales, residual-risk disclosures, and any PLSS posture declaration, as the downstream output consumed by ripple.md assurance records; (iii) a version-pin declaration listing the exact ripple.md and RippleLogic artifact versions and hashes used; and (iv) a disagreement-handling rule: if ripple.md feasibility gates and RippleLogic admissibility gates produce contradictory signals, the deployment MUST preserve both records, flag the disagreement, and escalate to governance review rather than silently collapsing the discrepancy.

Scope note (Informative). This interoperability section defines the portable interface contract between the two artifacts. It does not reproduce ripple.md internal semantics, which remain governed by the ripple.md Standard v5.8 canonical artifact. For the full assurance-case architecture, Decision Note schema, waiver choreography, and emergency triage semantics, see ripple.md v5.8.

End Appendix S.

Appendix T: Qualification Continuity for Consequence-Bearing Action (Normative for Tier 2-3 execution)

T.1 Purpose and architectural placement

Qualification continuity prevents a run from treating a prior calculation, selection, or approval as authority for a materially different action. It is an execution-bound conformance discipline applied after the existing RG -> RF/NCRC -> TRC -> CSV -> RLS cascade and before consequence-bearing action. It is not a sixth gate, a welfare dimension, a legal conclusion, or physical-safety proof.

T.2 Exact binding rule

For Tier 2 and Tier 3 consequence-bearing execution, the PCC or linked run record SHALL bind the selected option to all of the following at the time of action: action_qualification_id; qualified_option_id; action_instance_id; action_specification_hash; configuration_id and, where available, configuration_hash; qualification_snapshot_hash; gate-result snapshot; current-state status; transition status; post-state status; qualification status; active control and material-obligation status; monitoring plan; pause or shutdown path; rollback or repair path; separate authority-basis reference; reviewer identity or accountable role; evaluation time; and claim boundary.

The action specification SHALL identify the material action parameters, tools, permissions, operating envelope, affected scopes, and execution boundary. A change that would alter the applicable evidence, rights analysis, tail-risk analysis, containment analysis, control performance, authority basis, or affected parties requires a new action specification and requalification.

T.3 Execution-bound consistency rule

An execution-bound claim is conforming only when all of the following hold simultaneously:

1. `qualified_option_id` equals the currently selected selectable option;
2. current-state status is **VALID**;
3. transition status is **ADMISSIBLE** under the existing gates;
4. post-state status is **VIABLE** within the declared claim and evidence boundary;
5. qualification status is **CURRENT** and no material requalification trigger is active;
6. the exact action and exact configuration match the frozen qualification snapshot;
7. required controls and material obligations are current and effective for their declared role;
8. monitoring, pause or shutdown, and rollback or repair paths are usable within the consequence tempo; and
9. a separately legitimate authority basis applies to the same action, configuration, scope, and time.

Failure of any condition SHALL block ordinary execution and route the run to repair, narrower scope, escalation, or Refusal as the applicable existing rule requires. This rule cannot rescue an RF/NCRC, TRC, CSV, evidence, control, or authority failure.

T.4 Capability, selection, and authority separation

Capability availability, capability enablement, framework selectability, framework selection, execution qualification, and legitimate authority are distinct states. No state implies the next. MathGov supplies a bounded governance record; it does not itself grant institutional, democratic, legal, operational, or physical authority.

T.5 Material-unknown action rule

Where the validity domain, safe operating envelope, causal mapping, or failure boundary is materially unknown, governance MAY qualify only a bounded evidence-gaining action whose own risks pass RF/NCRC, TRC, and CSV and whose stopping, containment, monitoring, and review rules are predeclared. Governance SHALL NOT translate the unknown into implied safety or ordinary execution eligibility through approval, documentation, user acceptance, certification, or risk transfer.

T.6 Validation rule

For a required run, absence, mismatch, expiry, or trigger-suspension of the action-bound qualification SHALL make any execution-qualified claim **INVALID**. It does not rewrite the historical gate arithmetic; it limits what can be claimed or done from that arithmetic.

Appendix U: Outcome Observation and Requalification (Normative for consequence-bearing action)

U.1 Purpose

MathGov remains answerable to observed consequences. A completed action does not retroactively prove its assumptions, causal model, safety, or moral correctness. Observation updates the evidence record and reopens affected stages when the new evidence is material.

U.2 Required outcome observation

For Tier 2 and Tier 3 consequence-bearing action, the PCC or linked record SHALL identify: `outcome_observation_id`; `action_instance_id`; `qualified_option_id`; `configuration_id`; observation time and window; predeclared expected markers; observed outcomes; deviation from expected results; materiality assessment; affected Union Scopes and worst-affected subgroups; rights-floor, tail-risk, containment, control, and material-obligation signals; new material unknowns; source and measurement method; uncertainty; reviewer; and claim effect.

U.3 Requalification triggers

A run SHALL suspend inherited execution qualification and reopen the affected stages when observation reveals a material configuration change, new affected party, credible rights-floor concern, new ruin pathway, containment degradation, control or obligation failure, authority change, material deviation from predeclared outcomes, incident, near miss, cumulative harm, or a new unknown that changes the attempted claim.

Reopening is dependency-localized: the run SHALL reopen every stage and downstream result affected by the changed premise or evidence, but it need not discard unaffected evidence without reason. If the dependency boundary is itself underdetermined, the run SHALL use the more conservative reopening scope or enter Refusal.

U.4 No post-hoc rescue

Outcome evidence SHALL be assessed against predeclared markers, tolerances, and evidence rules. Implementations MUST NOT redefine success after observing the result, collapse subgroup harms into aggregate benefit, treat absence of observed failure as proof of safety, or preserve a preferred selection by silently narrowing the observation window.

U.5 Claim effect

Observation MAY support continuation, controlled continuation, redesign, suspension, rollback, repair, escalation, or Refusal. It SHALL NOT create authority. A successful bounded trial supports only the claim justified by its configuration, population, duration, measurement quality, and validity domain.

U.6 Machine-readable implementation

Run-record schema v4 and its semantic validator operationalize the execution-bound subset of qualification continuity: exact option/action/snapshot binding, current qualification, valid current state, admissible transition, viable post-state, and capability/execution-state consistency. Outcome requalification remains partly reviewer-governed because materiality and causal interpretation require domain evidence that a generic schema cannot establish.

APPENDIX V: Measurement Maturity and Deployment Context Hardening

Normative for Tier 2-3 claim-bearing runs; recommended for high-stakes Tier 1 runs.

V.1 Purpose

RippleLogic gates may be architecturally binding before their measurement surfaces are empirically mature. Every run MUST disclose the maturity posture of the evidence, indicators, kernels, institutional weighting process, and escalation mechanisms used to support those gates. A binary gate verdict is only claimable when the declared evidence basis can support the required determination under the run's declared thresholds and uncertainty rules.

V.2 UCI Measurement Readiness Classification

Token	Meaning	Claim effect
UCI-M0_UNAVAILABLE	No independent structural indicators are available.	UCI cannot support Containment or tie-break claims without escalation or narrower claim scope.
UCI-M1_PROVISIONAL	Indicators are declared but not validated, normalized, or reliability-tested.	Tier run may proceed only with disclosure and caution; no cross-domain UCI validation claim.
UCI-M2_DOMAIN_SUPPORTED	Instruments, normalization rules, and reviewer procedures are declared for the domain.	Stronger Tier 3 use is possible with disclosed provisionality.
UCI-M3_VALIDATED	A domain annex exists with reliability, construct-validity, normalization, and structural-independence evidence.	Higher-confidence Containment and tie-break use is claimable within that domain.

For Tier 2-3 runs using UCI, PCC.UCI_Measurement_Readiness, PCC.UCI_Indicator_Source, PCC.UCI_Normalization_Rule, PCC.UCI_Structural_Independence_Posture, and PCC.UCI_Reviewer_Reliability_Status SHALL be recorded. If UCI-M0 applies, UCI-dependent Containment or tie-break claims SHALL route to escalation, conservative fallback, or Refusal rather than silently proceeding as validated measurement.

V.3 Kernel Use Posture

Token	Meaning
K0_NONE	No numerical propagation kernel is used; direct multi-scope assessment and conservative disclosure govern.
K1_STARTER	Starter kernel used with shrinkage, sensitivity checks, and low confidence.
K2_DOMAIN_CALIBRATED	Domain-specific kernel has documented sources and partial validation.
K3_BACKTESTED	Kernel has backtest evidence, prediction-error documentation, and known failure boundaries.
K4_REGISTERED	Kernel is published, hash-pinned, independently reviewed, and tied to replayable test vectors.

PCC.Kernel_Use_Posture, PCC.Kernel_Validation_Evidence, PCC.Kernel_Backtest_Status, and PCC.Kernel_Error_Bounds SHALL be disclosed for any run using numerical propagation. K0_NONE is conformant when disclosed; it means the run is ripple-aware through scope mapping, gates, and auditability, but does not claim validated numerical propagation.

V.4 HDW Capture Audit Record

For Tier 3 public or institutional decisions, HDW weight-setting SHALL include a capture-risk assessment. Required fields: PCC.HDW_Weight_Source, PCC.HDW_Stakeholder_Map, PCC.HDW_Participation_Quality, PCC.HDW_Expert_Role_Provenance, PCC.HDW_Conflict_Disclosures, PCC.HDW_Challenger_Window, PCC.HDW_Independent_Review_Status, and PCC.HDW_Captured_Context_Limitation. Transparent capture is still capture; disclosure does not repair captured authority.

V.5 PLSS Escalation Trigger Register

PLSS runs SHALL record whether hidden externalities, rights exposure, precedent effects, lock-in, ecological spillover, governance spillover, repeated local use, or non-local stakeholder exposure are plausible. Required fields: PCC.PLSS_Profile_Used, PCC.Local_Stakeholder_Set, PCC.Guardrail_Scopes, PCC.Escalation_Trigger_Checklist, PCC.Non_Escalation_Rationale, PCC.Externality_Search_Method, PCC.Prominence_Sensitivity_Result, PCC.External_Challenger_Status, and

PCC.Repeated_PLSS_Use_Flag. Tier 2 PLSS runs require prominence-sensitivity review when non-local effects are plausible; Tier 3 PLSS runs require external challenger review or documented waiver.

V.6 Deployment Context Profile

Every Tier 2-3 claim-bearing deployment SHOULD classify its deployment context as LAB, PILOT, INSTITUTIONAL_ADVISORY, INSTITUTIONAL_OPERATIONAL, PUBLIC_POLICY, AGENTIC_EXECUTION, or EMERGENCY_TRIAGE. Context determines evidence sufficiency, reviewer independence, audit visibility, and escalation thresholds.

V.7 Non-Overclaiming Rule

A run MAY be architecturally valid while measurement-maturity is provisional. It MUST NOT claim empirical validation, cross-domain calibrated UCI, validated kernel propagation, biological SGP measurement completeness, capture-resistant HDW, or safe PLSS localism unless the relevant maturity fields support that claim.

V.8 Machine-Verifiability and Calibration-Prep Addendum

Cumulative Release Delta (v10.2 → v10.4; lineage context). This cumulative lineage block summarizes changes accumulated across the v10.2-to-v10.4 development arc. It is not the immediate point-release characterization of v10.4, which is defined in Appendix M.4A and the Methodology Upgrade Summary as a readability, front-matter reorganization, and companion synchronization release. The cumulative arc preserves the canonical admissibility cascade, equations, rights thresholds, TRC mechanics, Containment semantics, SGP binding interface, PLSS/FLGG architecture, and Tier 1-3 claim boundaries, while adding schema ownership rules, registry versioning rules, machine-readable registry expectations, UCI Measurement Annex linkage, PLSS calibration linkage, Appendix R machine-vector requirements, reference-calculator expectations, and CI drift-check requirements.

Normative effect statement (carried forward from v10.8). No intended change was made to the canonical admissibility cascade Reality Grounding → Rights Floor → TRC → CSV → RLS, rights coverage sets, TRC corridor mechanics, Containment semantics, RLS formula, PLSS/FLGG architecture, HDW, PCC, RSP/RefStructRecord, or Appendix AD. The v10.8 line updated the SGP companion pin to SGP v5.6 and clarified the human plateau / reality-management rationale without allowing any SGP protection surface to weaken human protections, alter admissibility gates, or grant governance authority by itself. The current v13.0 / SGP v8.8 line carries forward replacement of the legacy scalar interface with MPS/FPP/GPR/SPR/ICP and preserves the informative RMCP/P100 calibration surface while preserving protection, participation, capacity, and authority boundaries.

RippleLogic Canon v13.0 is the controlling specification text for this delivered Core 15 research and teaching package. Its readiness is bounded by the current verification report and unresolved evidence. It does not establish empirical validation, ProofPack readiness, complete machine-ecosystem conformance, legal certification or deployment certification.

V.9 Schema Ownership and Machine-Readable Source of Truth

For v11.0 and later machine-verifiable surfaces, each schema, registry, record, enum, and validation object SHALL have exactly one owning artifact. The owner artifact defines the object's fields, allowed values, semantics, and versioning rules. Consumer artifacts MAY reference the object by id, version, and manifest hash, but MUST NOT independently redefine it.

On conflict, the owner artifact controls. Consumer artifacts SHALL cite, import, or reference the owner's machine-readable registry or schema instead of duplicating field definitions.

The schema ownership map is the controlling coordination layer for current-line machine-verifiable work. It does not alter the Cascade, NCRC, TRC, CSV, RLS, UCI/HOI, SGP binding interface, or lawful authority boundaries.

V.10 UCI Measurement Annex Linkage

Where a UCI Measurement Annex is available and hash-pinned, implementations SHOULD cite the annex version, indicator registry, normalization rule, staleness rule, missing-data rule, and reviewer protocol used. In the absence of such an annex, UCI claims SHALL remain at the declared measurement-readiness level and MUST NOT be described as cross-domain validated.

For this v13.0 release line, UCI Measurement Annex v0.1 is provisional and initially scoped to U3 Community, U4 Organization, and U5 Polity contexts. It does not establish cross-domain validated UCI status.

V.11 PLSS Calibration and Adjudication Linkage

Tier 3 PLSS use SHOULD compare the declared `q_u` profile against a hash-pinned PLSS calibration or adjudication pack where available. Reviewer disagreement, escalation-trigger disputes, accepted or rejected local-scope rationales, externality-search results, and non-escalation rationales SHALL be recorded in the PCC.

The calibration pack does not create empirical validation by itself. It supports reviewer convergence, interpretive consistency, and anti-capture discipline.

V.12 Machine-Readable Test Vector Rule

For any release that claims machine-verifiable, ProofPack-prep, reference-calculator, or validator-supported status, Appendix R SHALL be accompanied by a machine-readable companion file. The current v13.0 core source release does not bundle such files unless they are explicitly listed and hash-pinned in the release manifest. In the absence of manifest-listed machine-readable vectors, Appendix R remains the human-readable normative reference surface, and no machine-verifiable ecosystem-completeness claim is made. When bundled, the companion SHOULD include input options, union and dimension weights, rights impacts, TRC scenario inputs, Containment inputs, RLS expected outputs, UCI/HOI diagnostic inputs where applicable, expected gate verdicts, expected selected or refusal disposition, tolerance rules, and expected error behavior for failure-path examples.

A human-readable Appendix R remains normative for explanatory purposes in the current v13.0 core release. Where a machine-readable companion is explicitly manifest-listed for a future release or tooling bundle, that companion becomes the required reference surface for validators and reference calculators within the claim scope declared by that release.

V.13 False Precision Guardrail (Normative for Tier 2-3 claim-bearing runs)

A numerical cell value is not evidence by itself. Every high-impact, high-uncertainty, gate-sensitive, or culturally contested cell SHOULD disclose its evidence source, uncertainty class, reviewer confidence, and sensitivity relevance. Unsupported numerical precision in welfare impacts, probabilities, UCI indicators, SGP-consumed scalars, or PLSS scope declarations SHALL be treated as an epistemic risk rather than as objective certainty.

Where a cell value is weakly evidenced, likely to swing an NCRC/TRC/CSV verdict, likely to alter escalation or refusal status, or dependent on disputed subgroup or threshold definitions, the run SHOULD trigger epistemic-humility review, sensitivity analysis, reviewer-disagreement logging, escalation, narrower claim scope, or Refusal. Implementations SHOULD flag `Boundary_Gaming_Risk`, `False_Precision_Risk`, `Gate_Swing_Cell_Flag`, `Subgroup_Definition_Risk`, `Scenario_Probability_Risk`, `Threshold_Laundering_Risk`, and `Reviewer_Disagreement_Level` where applicable.

Meaning, agency, social cohesion, dignity, and ecological integrity cells require particular caution. These dimensions may be measurable in some contexts, but unsupported precision in these cells SHALL NOT be presented as stronger than the underlying evidence, testimony, instruments, or qualitative-to-quantitative coding rules allow.

APPENDIX W: SCHEMA OWNERSHIP, REGISTRIES, AND MACHINE-VERIFIABLE SURFACES

W.1 Purpose

This appendix defines ownership boundaries for machine-readable schemas, registries, generated tables, and validator surfaces. It prevents duplicate schema authorship, manual synchronization drift, and hidden field-definition conflicts across Canon, SGP, ripple.md, Agent System, workbook, and release-integrity files.

W.2 Ownership Rule

Each schema or registry has exactly one owner. The owner defines the object. All other artifacts consume the object by id, version, and manifest hash.

W.3 Canon-Owned Objects

The Canon owns PCC core record, OutcomeObservationRecord, ActionQualificationRecord, RippleLogic PrimitiveTraceRecord, FalsificationPlan, ObservableConsequenceRegister, ProjectionPreRegistration, MaturityProfileRecord, Appendix R test vectors, rights registry, catastrophe-cell registry, reason-code registry, uncertainty-method registry, PLSS operator registry, deployment-context enum registry, and Canon audit flags.

W.4 SGP-Owned Objects

SGP owns the SGP evaluation record, MPS evidence-family fields, MPS band/interval and MPS-NE status, FPP/GPR/SPR/ICP typed fields, target and identity records, evidence-tier classifications, biological evaluation maturity tokens, and moral-patienthood evidence trace.

W.5 ripple.md-Owned Objects

ripple.md owns Decision Note schema, Assurance Preconditions Record, waiver records, triage records, update-legitimacy records, wrapper-level refusal records, and wrapper assurance records.

W.6 Agent-System-Owned Objects

The Agent System owns operator envelope, runtime audit log, mode-control records, rollback records, operator authentication records, and runtime conformance records.

W.7 Split Objects

Primitive trace is split: RippleLogic PrimitiveTraceRecord belongs to the Canon; SGP moral-status evidence trace belongs to SGP. Maturity tokens are split: UCI Measurement Readiness, Kernel Use Posture, HDW Capture Audit, PLSS Escalation, and Deployment Context tokens belong to the Canon; Biological Evaluation Maturity tokens belong to SGP. Audit flags are split: Canon flags affect admissibility, claimability, escalation, refusal, or conformance; tool-local flags affect workbook or implementation review surfaces only and do not by themselves invalidate a run unless they trigger a Canon flag.

W.8 Versioning Rule

External implementation-name crosswalk (Normative clarification). The Reproducibility Standard's schemas/mathgov_run_record_v4_1.schema.json, docs/implementation/NORMATIVE_KERNEL_INDEX_v1.1.yaml and release/VALIDATE_MATHGOV_RUN.py identify required external implementation surfaces. The Cascade Standard's CANONICAL_STATE_REGISTRY_v1.1.yaml and STATE_TRANSITION_MATRIX_v1.1.json identify the external token and transition surfaces under Canon ownership. Together with CANONICAL_AUDIT_FLAG_REGISTRY_v1.1.yaml, none becomes an available or verified package authority

merely by being named. The manifest and `Publication/Dependency_Status.json` distinguish bundled artifacts from unavailable dependencies. Until the exact interface files, hash pins, compatibility record and required tests are supplied, claims dependent on those external interfaces remain unavailable; the controlling prose and the documented supplemental subset remain distinct.

Machine-readable registries and schemas **SHOULD** use stable filenames with an internal version field. The release manifest **SHALL** hash-pin each registry and schema. The manifest hash, not filename churn, controls byte-level release identity.

W.9 Generator Rule

Generated documentation tables, workbook validation ranges, verifier constants, README companion matrices, and manifest scaffolds **SHOULD** be generated from machine-readable registries whenever possible. Generated files **SHOULD** include a generated-from note and **SHOULD NOT** be manually edited except through the owning registry or schema.

W.10 ProofPack Boundary

The existence of schemas or registries does not by itself create ProofPack readiness. ProofPack readiness requires schemas, validator or CLI, reference calculator, canonical test vectors, replayable PCC bundle, independent replay instructions, hash-pinned replay materials, and public third-party replay path.

APPENDIX X: REFERENCE CALCULATOR, MACHINE VECTORS, AND PROOFPACK-PREP ROADMAP

X.1 Purpose

This appendix records the current-release ProofPack-prep path for machine-readable Appendix R test vectors, starter reference-calculator behavior, golden outputs, validator linkage, and independent replay readiness. It does not create ProofPack readiness by itself.

X.2 Machine-readable Appendix R vectors

Future ProofPack-prep and tooling companion releases **SHOULD** include machine-readable Appendix R vectors. This candidate bundles the bounded regression evidence listed in its manifest; that evidence is not a complete machine-readable Appendix R conformance surface. When bundled, vectors **SHOULD** include inputs, union and dimension weights, rights impacts, TRC scenario inputs, Containment inputs, RLS expected outputs, UCI/HOI diagnostic inputs where applicable, expected gate verdicts, selected or refusal disposition, tolerance rules, and expected error behavior for failure-path examples.

X.3 Reference Calculator Starter

The reference calculator starter is a scaffolding artifact for reproducing Appendix R expected outputs. It is not a full semantic validator, not a production-grade implementation, and not empirical validation. It becomes claim-strengthening only when it reproduces golden outputs under declared tolerance rules and is independently replayable.

X.3A Reference Calculator Priority (ProofPack-prep guidance)

The next implementation priority is an open reference calculator that reproduces the Canon Appendix B equations and Appendix R test vectors. Reference implementations **SHOULD** match golden gate verdicts exactly and numerical outputs within declared tolerance rules unless exact arithmetic is used.

The reference calculator **SHALL NOT** be described as empirical validation, a full semantic validator, or ProofPack completion until it is paired with schemas, replayable PCC artifacts, independent replay instructions, manifest/hash pinning, and public third-party replay. Its immediate purpose is narrower and

concrete: to make Tier 1-3 calculation reproducible enough for reviewers to detect formula drift, implementation disagreement, and invalid claim expansion.

X.4 CI and generated-consistency checks

CI checks SHOULD verify manifest hashes, registry/schema ownership, stale pins, generated-file drift, workbook public surfaces, and external-reference sanitation. These checks improve release hygiene but do not establish empirical validation or Tier 4 readiness.

X.5 ProofPack boundary

ProofPack readiness requires schemas, validator or CLI, reference calculator, canonical test vectors, replayable PCC bundle, independent replay instructions, hash-pinned replay materials, and a public third-party replay path. This v13.0 specification claims neither ProofPack-prep completeness nor ProofPack readiness.

Appendix Y: Precision Patch - Legitimacy, Evidence Status, Validation, and Adoption Profiles

Y.1 Release Delta (v10.3.1 → v10.4)

This patch preserves the v10.3.1 cascade architecture, equations, rights thresholds, TRC mechanics, Containment semantics, SGP binding interface, PLSS/FLGG architecture, and Tier 1-3 claim boundaries. It adds precision safeguards for public legitimacy, evidence-status labeling, SGP patienthood/participation separation, compact institutional profiles, and the validation roadmap. No empirical-validation, ProofPack, Tier 4, machine-verifiable ecosystem-complete, or reference-calculator-complete claim is created by this patch.

Y.2 Public Legitimacy Boundary (Normative Clarification)

RippleLogic may structure, audit, and improve public reasoning, but it does not itself confer democratic legitimacy, legal authority, public mandate, professional licensing, or institutional permission. In public governance contexts, legitimacy requires lawful authority, affected-stakeholder participation, public reason, contestability, appeal or review procedures, and external accountability outside the framework itself. A RippleLogic run may support public justification; it MUST NOT be presented as a substitute for democratic authorization, legal compliance, or domain-specific professional judgment.

Y.3 Evidence Status Tokens (Normative for Tier 2-3 claim-bearing records)

For Tier 2-3 claim-bearing records, decision-critical numerical or categorical inputs SHOULD carry an evidence-status token. Where the token is unavailable, the field SHOULD default to UNVALIDATED or NOT_OBSERVED rather than silently appearing as measured fact.

Y.4 Required Evidence-Status Token Registry

MEASURED: directly observed or measured with disclosed method and source. ESTIMATED: calculated or inferred from disclosed evidence. ASSUMED: explicit assumption adopted for the run. EXPERT_JUDGMENT: reviewer or domain expert judgment with role and basis recorded. CONTESTED: materially disputed evidence or interpretation. UNVALIDATED: plausible but not externally validated for the claimed use.

NOT_OBSERVED: no direct evidence currently observed. PLACEHOLDER_FOR_SENSITIVITY_ONLY: value used only for sensitivity exploration and not as a factual claim.

Y.5 Compact Institutional Profiles (Informative Roadmap)

Implementations MAY publish compact profiles for specific institutional contexts, provided they do not waive NCRC, TRC, CSV, Refusal, SGP boundary protections, public legitimacy requirements, or claim-boundary disclosures. Initial compact profiles SHOULD include: RippleLogic-enhanced RIA Profile, RippleLogic

Procurement/Vendor-Risk Profile, RippleLogic Agent Deployment Gate Profile, and RippleLogic University AI Policy Review Profile. Each profile should state required inputs, mandatory gates, simplifications allowed, non-waivable controls, output artifact, and escalation triggers.

Y.6 Validation Roadmap Commitments (Informative, non-validation claim)

The next validation layer SHOULD prioritize: reference-calculator fidelity; JSON/YAML schemas for PCC, Decision Note, TRC, CSV, RLS, and SGP interface objects; canonical test vectors; replayable PCC examples; inter-rater reliability for NCRC/TRC/CSV/RLS/SGP; reconstructability audits; PLSS boundary-gaming red-team studies; HDW legitimacy pilots; UCI measurement annex pilots; SGP biological evaluation annex development; and shadow-mode pilots in regulatory impact analysis, public procurement, university AI policy, organizational ethics review, and AI agent deployment review.

Appendix Z: Integration-Hardening Addendum

Normative effect statement. This addendum does not alter the canonical cascade, rights thresholds, TRC corridor mechanics, SGP MPS/FPP binding, PLSS/FLGG architecture, or Tier 1-3 claim boundaries. It is a precision integration patch that tightens workbook linkage, evidence-status disclosure, SGP split disclosure, and release-verification discipline.

Z.1 Workbook integration rule

Any workbook companion that adds claim-boundary, evidence-status, SGP-split, validation-roadmap, or compact-profile sheets SHOULD include active sanity checks confirming that those sheets are present, populated, and referenced by at least one workbook validation surface. A sheet used to support a release claim MUST NOT remain merely decorative.

Z.2 Evidence-status propagation rule

Decision-critical inputs SHOULD disclose an Evidence Status Token. Where a decision-critical input lacks a token, the PCC SHOULD record NOT_OBSERVED, UNVALIDATED, or ASSUMED rather than silently treating the input as measured evidence.

Z.3 SGP split computation bridge

MPS, GPR, and SPR may be computed or recorded in workbook companions as disclosure outputs. MPS concerns protection-relevant welfare-bearing standing. GPR concerns role participation readiness. SPR concerns responsibility-bearing authority. Low GPR or SPR MUST NOT reduce MPS or weaken any NCRC, TRC, or CSV protection.

Z.4 Verifier sufficiency rule

A release verifier SHOULD test more than filename presence. For workbook-linked additions, it SHOULD check required sheet names, required labels, non-empty expected fields, and representative formula references. This does not make the release ProofPack-ready or machine-verifiable ecosystem complete; it only reduces packaging and integration risk.

Z.5 Claim boundary

This addendum improves release integrity and companion calculability. It does not claim empirical validation, ProofPack readiness, Tier 4 readiness, complete reference-calculator status, or machine-verifiable ecosystem completeness.

APPENDIX AA: PACKAGE-CONTENT BOUNDARY AND TOOLING LINE SEPARATION (CARRIED-FORWARD ORIGIN; ACTIVE REQUIREMENTS)

AA.1 Purpose

This appendix clarifies the boundary between the current v13.0 core source package and companion tooling releases. It is added to prevent asset-presence overclaiming and to keep the canonical framework synchronized with the actual release manifest.

AA.2 Package-content rule

Where this canon references schemas, registries, machine-readable vectors, annexes, reference calculators, CI workflows, evidence registries, validators, or replay bundles, those references describe ownership rules, interface expectations, or roadmap targets unless the corresponding artifact is explicitly listed in the release manifest for the package being claimed.

AA.3 Tooling-line separation

The package tooling line is the proper home for schemas, reference-calculator modules, machine-readable test vectors, replayable PCC examples, validators, evidence-hash registries, and external audit protocols. The current v13.0 core source line is the controlled conceptual and normative candidate; the v12.7 line is preserved as the immediate audited baseline and legacy-replay source. Tooling artifacts may support replayability, but they do not supersede the canon.

AA.4 Claim boundary

The existence of tooling, schemas, test vectors, or validators does not by itself create ProofPack readiness, Tier 4 readiness, empirical validation, machine-verifiable ecosystem completeness, legal authority, democratic legitimacy, or consciousness detection. Stronger claims require the separate evidence and replay conditions stated in the relevant release contract.

APPENDIX AB: REFERENCE STRUCTURE PROTOCOL

AB.1 Purpose

The Reference Structure Protocol specifies how RippleLogic runs declare the structures against which admissibility, causality, stability, welfare, and auditability claims are evaluated.

This appendix prevents a common failure mode in governance and AI alignment: one assumption system supervising another assumption system without a sufficiently external, falsifiable, or auditable reference. RippleLogic does not solve this by claiming perfect objectivity. It solves it by requiring that reference structures be declared, evidence-labeled, maturity-bounded, and open to refusal.

AB.2 Non-Effect on the Cascade

The Reference Structure Protocol does not alter the canonical cascade: Reality Grounding → Rights Floor → TRC → CSV → RLS.

It does not create a sixth gate. It does not allow RLS, UCI, HOI, SGP, agent tooling, legal permission, or human authority to rescue failure at NCRC, TRC, CSV, or Refusal.

AB.3 Reference Stack

A RippleLogic run SHOULD identify which reference layers are active.

1. Physical reference: bodies, infrastructure, energy, geography, resources, physical constraints, ecological limits.
2. Biological reference: life, health, basic needs, suffering, sentience-relevant evidence, biological dependency.
3. Social reference: trust, coordination, institutional stability, conflict risk, social cohesion.
4. Ethical-rights reference: rights floors, dignity, due process, non-domination, ecological integrity.
5. Risk reference: catastrophic corridors, CVaR assumptions, scenario sets, probability floors, uncertainty bands.
6. Governance-authority reference: lawful authority, role authority, democratic legitimacy, professional mandate, procedural permission.
7. Epistemic reference: evidence source, measurement method, evidence status, falsification condition, uncertainty class, reviewer disagreement.
8. Operational-audit reference: PCC, Decision Note, ProjectionPreRegistration, ActionQualificationRecord, OutcomeObservationRecord, hashes, logs, and replay materials.

Short rule. Physics constrains possibility; ethics constrains admissibility; evidence constrains claims; governance constrains authority.

AB.4 ReferenceStructureRecord

For Tier 3 runs and for Tier 2 high-stakes or claim-bearing runs, the PCC SHOULD include a ReferenceStructureRecord. For Tier 3 public conformance claims, it SHALL include one unless an emergency triage rule explicitly permits a narrower record.

A ReferenceStructureRecord contains, at minimum: record_id; decision_id; option_id or option_set_id; claim_scope; decision_boundary; reference_object; active_reference_layers; invariant_constraints; causal_pathway_summary; falsification_conditions; evidence_sources; evidence_status_tokens; measurement_maturity_status; adaptation_boundary; authority_boundary; refusal_triggers; excluded_references; reviewer_disposition; timestamp; artifact_hashes.

AB.5 Reference Quality States

Each major reference claim SHOULD be classified using one of the following states: MEASURED; ESTIMATED; EXPERT_JUDGMENT; ASSUMED; CONTESTED; UNVALIDATED; NOT_OBSERVED; PLACEHOLDER_FOR_SENSITIVITY_ONLY.

Assumed, contested, unvalidated, not-observed, or placeholder references MUST NOT be presented as measured facts.

AB.6 Invalid or Insufficient Reference Structures

A reference structure is insufficient for a strong RippleLogic claim where any of the following conditions hold:

1. The reference object is missing.
2. The reference is circular or self-validating.
3. The reference is controlled entirely by the system being evaluated without independent evidence or review.
4. The reference cannot be connected to an observable consequence, threshold, or falsification condition.
5. The reference is materially stale for the decision horizon.
6. The reference depends on a measurement surface whose maturity status does not support the requested claim.

7. The reference relies on legal, institutional, or operator permission as a substitute for admissibility.
8. The reference relies on scoring outputs to rescue earlier gate uncertainty or failure.
9. Two or more materially different reference interpretations remain plausible under the declared evidence and no adjudication rule resolves them.

Where the reference structure is insufficient, the run **MUST** enter Refusal, escalation, narrower claim scope, additional evidence gathering, or bounded triage if delay itself creates material risk.

AB.7 Relationship to PCC and ProjectionPreRegistration

The ReferenceStructureRecord **SHOULD** be stored with the PCC and referenced by hash. Where ProjectionPreRegistration is required, the ReferenceStructureRecord **SHOULD** be created before cascade computation or included as a pre-registered object.

Post-hoc modification of the reference structure that could alter admissibility, selectability, selection, conformance, or deployment-readiness claims is material and requires a new run or documented refusal and re-run procedure.

APPENDIX AC: SPECIFICATION CONTRACT, RELEASE METADATA, AND CLAIM BOUNDARY

This appendix preserves release identity, companion-version pins, claim-boundary warnings, and role-map material moved out of the opening pages in v10.4 for reader accessibility. It does not change the canonical cascade or computation.

Foundation baseline consolidation date: 2026-05-11. Canonical release date: governed by the release manifest and the current-line release entry in Appendix M. If dates differ, the release manifest date controls release identity; the Appendix RELEASE preparation and revision dates record document finalisation; the exact build manifest and hashes distinguish corrected deliveries.

Release Label: RippleLogic_v13.0_Canon

Current synced companion set for this release line: RippleLogic Agent System v13.0; SGP v8.8; ripple.md v5.8; WDBIP v1.9; RippleLogic Aligners Sheet v5.9; Foundations Primer v4.7; Cascade Standard v2.9; CSV Gate Standard v2.7; Reproducibility and Use Standard v1.7; RLS Validation Protocol v2.9; RLS Validation Workbook v0.4 (release status governed by repository manifest and SHA-256 hashes).

Delivered release-registry rule. VERSION_MANIFEST.json and SHA256SUMS.txt identify this package's exact current component bytes. VERSION_MANIFEST.yaml and release/release_manifest.yml remain external repository authorities and require separate synchronization before any repository-conformance claim. A mismatch inside this delivery is a release-integrity failure; publication status remains separate from the component edition.

Current-line companion matrix (release-clean reference):

- RippleLogic Canon: v13.0
- SGP: v8.8
- ripple.md: v5.8
- WDBIP: v1.9
- RippleLogic Agent System: v13.0
- Foundations Primer: v4.7
- RippleLogic Aligners Sheet: v5.9
- Cascade Standard: v2.9
- CSV Gate Standard: v2.7
- Reproducibility and Use Standard: v1.7

- RLS Validation Protocol: v2.9
- Public Introduction: v13.0
- PC-AEP: v2.6
- MFDI: v2.6
- Source-Coupling Integrity Standard: v2.6

Extended release-support artifact (not part of Core 15):

- RLS Validation Workbook: v0.4

All public conformance claims for this release line **MUST** pin this exact component matrix explicitly in the release manifest or PCC; generic release-family references do not replace the exact current component pins and the exact manifest controls publication status.

MathGov is the umbrella ethical governance framework. RippleLogic is the canonical name of the decision architecture and operating methodology inside MathGov. Earlier materials may use historical lineage language, but current release interpretation follows this hierarchy.

Backward compatibility: Prior PCCs computed under earlier release lines remain interpretable as lineage artifacts. However, any run claiming current-line conformance **MUST** meet v13.0 canonical requirements and use the v13.0 Canon artifact. A run explicitly pinned to v12.0 remains a historical-lineage run and must pin the v12.0 artifact. Any run referencing v10.8 remains a historical-lineage run and must pin the v10.8 artifact explicitly.

AUTHOR STATEMENT AND TRANSPARENCY DISCLOSURE

The author conceived, designed, and wrote this specification. Generative AI tools (including OpenAI ChatGPT and Anthropic Claude) were used as drafting and consistency assistants. The author reviewed, verified, and edited all outputs and assumes full responsibility for the content, claims, and citations. This paper presents a theoretical framework. Empirical validation through pilots is planned but not yet completed. No operational deployment claims are made until validation studies are completed. This authorship disclosure describes the baseline. Candidate corrections were prepared through AI-assisted implementation and review; author review or approval of this candidate is not presumed.

Maturity note. This candidate preserves the framework architecture and implements bounded audit corrections. Empirical validation, design-specific reliability criteria, companion machine surfaces, domain measurement annexes and remaining audit decisions are still open; architectural completeness or operational readiness is not inferred from this build.

Historical release-delta details for v9.6.4, v9.6.5, v10.0-RC1, v10.0-RC2, v10.0, and v10.4 are retained in Appendix M. Appendix RELEASE states the current package release posture.

CARRIED-FORWARD v10.8 HUMAN PLATEAU / CMIU CALIBRATION AND SGP v5.6 SYNCHRONIZATION NOTICE

Release Delta (v10.5 → v10.6). This release integrates the Human Plateau / CMIU calibration clarification and synchronizes the Canon to SGP v5.6. It preserves the existing cascade architecture, equations, rights thresholds, TRC mechanics, Containment semantics, RLS formula, PLSS/FLGG architecture, HDW, PCC, RSP/RefStructRecord, Appendix AD scoring dictionary, and Tier 1-3 claim boundaries. It clarifies that the Human Plateau Rule is not species favoritism: human persons receive non-degradable protection, and humanity is the current evidenced plateau-class reality-managing community within the CMIU, while equivalent plateau extension remains open through SGP governed review.

Normative effect statement (carried forward from v10.8). No intended change was made to the canonical admissibility cascade Reality Grounding → Rights Floor → TRC → CSV → RLS, rights coverage sets, TRC corridor mechanics, Containment semantics, RLS formula, PLSS/FLGG architecture, HDW, PCC, RSP/RefStructRecord, or Appendix AD. The v10.8 line updated the SGP companion pin to SGP v5.6 and clarified the human plateau / reality-management rationale without allowing any SGP protection surface to

weaken human protections, alter admissibility gates, or grant governance authority by itself. The current v13.0 / SGP v8.8 line carries forward replacement of the legacy scalar interface with MPS/FPP/GPR/SPR/ICP and preserves the informative RMCP/P100 calibration surface while preserving protection, participation, capacity, and authority boundaries.

Claim boundary table (Informative reading aid).

Release-boundary clarification (Normative).

“Architecturally complete,” “worked-run claimable,” “audit-ready,” “ProofPack-ready,” and “Tier 4-ready” are distinct statuses and MUST NOT be collapsed.

A framework may be architecturally complete without being empirically validated. A worked run may be claimable within its declared scope without the ecosystem being ProofPack-ready. Auditability surfaces may exist without a public replay validator. Tier 4 readiness remains prohibited until a public, independently replayable ProofPack exists.

Phase 1 release-boundary hardening (Normative clarification).

For public release interpretation, RippleLogic v13.0 is a conceptual/core-foundation and Tier 1-3 architectural specification within its declared scope. It is not an empirical validation claim, not a ProofPack-ready claim, not a Tier 4 claim, not a machine-verifiable-ecosystem-completeness claim, not a cross-domain validated UCI claim, not a legal-compliance certification, and not a deployment-readiness certification.

The phrase “ethical operating system” is a design-architecture claim. It SHALL NOT be read as a claim of universal cultural consent, lawful authority, empirical completion, automatic moral correctness, or permission to execute consequence-bearing decisions without the required governance authority and evidence.

UCI maturity boundary. Structural coherence and viability remain load-bearing; UCI is a governed but empirically provisional diagnostic and tie-break evidence interface until a published UCI Measurement Annex defines instruments, normalization rules, structural-independence checks, calibration procedure, and inter-rater reliability evidence. UCI has no independent pass/fail or selection authority merely because it is numerically computable. No current-release artifact may imply cross-domain validated UCI unless that annex and evidence are separately bundled and hash-pinned.

Admissibility vs permission (Normative clarification). Admissibility is a structural property determined by the cascade gates. Permission is a separate authorization act performed by an entity with appropriate authority. Compliance is a documentation property. A run that passes the cascade is admissible. An admissible run may still require permission, lawful authority, governance review, or domain-specific licensing before execution. A permitted run is not, by virtue of permission alone, admissible. RippleLogic determines admissibility; it does not grant permission, certify compliance, or substitute for lawful authority.

Claim type	Status in v13.0	Required condition
Tier 1-3 architectural completeness	Claimable	Defined by this canon and appendices for the declared release scope.
Tier 3 run-level conformance	Claimable	Run declares required inputs, satisfies tier minima, and emits a valid PCC.
Framework-level empirical operational-readiness	Not yet claimable	Requires empirical validation progress, including UCI measurement maturity and related evidence.
Machine-verifiable ecosystem completeness	Not yet claimable	Requires companion validator, schema, reference-calculator / registry surfaces, or equivalent ProofPack assets.
Tier 4 compliance	Prohibited	No Tier-4 claim until ProofPack is publicly replayable and independently verified.
Cross-domain calibrated or validated UCI	Prohibited	Requires a published UCI Measurement Annex with named instruments, normalization rules, structural-independence evidence, and IRR data.

Decision-state summary (Normative shorthand).

RippleLogic uses five distinct decision states that implementations MUST NOT collapse.

Rights-admissible = passes NCRC.

Admissible = option has RG_status in {RG_SUPPORTED, RG_NARROWED}, passes NCRC, and has TRC_status in {TRC_PASS, TRC_NOT_TRIGGERED}. The term does not imply execution authority.

Selectable = admissible and has CSV status in {CSV_PASS, CSV_PASS_WITH_CONTROLS, CSV_NOT_MATERIAL}.

Selected = chosen from the selectable set after RLS, tie-break, and any required escalation disposition.

Refusal = the framework cannot uniquely determine the applicable admissibility or selection state from declared inputs, evidence, projection rules, or layer commitments.

A later state MUST NOT be used to infer an earlier one, and no downstream score may rescue failure at an earlier state. Refusal MUST NOT be silently converted into gate failure, gate success, RLS tie-break, or discretionary selection. Refusal means the run lacks sufficient determination for the claim being attempted.

Qualification Continuity Check. Qualification continuity is not a sixth decision state and not a replacement gate. It is an execution-bound conformance surface for Tier 2 and Tier 3 consequence-bearing action, and recommended for medium- or high-stakes Tier 1 action. It verifies that the selected option, exact action, configuration, evidence, controls, obligations, and separate authority basis remain current before consequence becomes real. It cannot rescue a failure at NCRC, TRC, or CSV, cannot convert Refusal into permission, and cannot use a prior qualification to justify a materially changed action.

CORE ARTIFACT ROLE MAP (release-line orientation)

- RippleLogic_v12.7_Canon.md: prior governing repository source, retained as an external baseline reference. A synchronized v12.8 candidate repository source must be reviewed before promotion; it is not asserted to exist in this package.

- RippleLogic_v13.0_Canon.docx: independently hash-identified candidate text for controlled review. It preserves the baseline publication structure with logged changes; it does not claim parity with an absent candidate Markdown source or replace the published v12.7 source.

- SGP v8.8: companion moral-patienthood, protection, participation, capacity, and power-readiness protocol. It is authoritative for MPS, FPP, GPR, SPR, optional ICP, RMCP, open P100 calibration, target and identity/collective boundaries, evidence-dependence controls, precaution, governed review, and classification-stability cautions. RippleLogic consumes only typed SGP outputs through Appendix G and MUST NOT define or recompute SGP internals.

SGP split-roadmap boundary (Informative but release-controlling for interpretation).

SGP v8.8 separates moral patienthood, full protection, governance participation, stewardship/power readiness, intelligence, and reality-management capacity as first-class outputs. RippleLogic v13.0 consumes the pinned SGP v8.8 protocol interface; the separately versioned RMCP machine schema remains v8.4. RippleLogic MUST NOT infer sentience, intrinsic moral worth, governance authority, voting weight, office eligibility, CMIU membership, or stewardship power from RMCP/P100, MPS, or FPP.

This roadmap note does not change the RippleLogic v13.0 / SGP v8.8 binding interface. It prevents over-reading any MPS band or interval as a consciousness probability or cardinal moral-worth score, and prevents over-reading RMCP/P100 as species supremacy, wisdom, or governance authority.

- ripple.md v5.8: portable alignment contract and assurance wrapper for implementations and deployments, including the wrapped-run ActionQualificationRecord field for high-stakes consequence-bearing episodes.

- RippleLogic Agent System v13.0: deployment/control specification for agentic implementations under the current pinned companion line unless and until a later synchronized agent-system release is promoted.

- Ripple Aligners Sheet v5.9: worked-run exemplar companion for replay, disclosure, and training, including Qualification_Continuity and Outcome_Requalification sheets; not the repository-level validator or schema surface.
- Machine-verifiable ecosystem completeness requires additional companion surfaces when claimed, including validator / schema / reference-calculator / registry artifacts or ProofPack assets beyond this core canon line.
 - Companion v8.8 interface note: RippleLogic treats MPS band/interval, FPP, GPR, SPR, optional ICP, and RMCP/P100 as separate pinned outputs. MPS records control protection posture and the required welfare-inclusion hypothesis sensitivity; they are not cardinal welfare multipliers. FPP controls non-downgrade protection; GPR and SPR remain role-specific; RMCP/P100 remain informative capacity-calibration outputs. No SGP output may weaken admissibility under NCRC, TRC, or CSV, and no legacy v7 scalar may be silently translated.

AC.1 Specification Contract (moved from front matter in v10.4)

SECTION 0: SPECIFICATION CONTRACT (NORMATIVE)

0.1 Normative hierarchy (Single Source of Truth)

When interpreting RippleLogic v13.0, conflicts MUST be resolved using this precedence order (highest controls lowest):

1. Main text rules explicitly labeled MUST, SHALL, or PROHIBITED
2. Appendix B (Canonical Equations)
3. Appendix C (Rights / NCRC Canon Pack)
4. Appendix D (TRC and Scenario Governance Canon Pack)
5. Appendix H (PCC and Audit Flags Canon Pack)
6. Appendix AG (v10.8 Reality Grounding and Lean Cascade Patch)
7. Appendix AH (Category Grounding and Primitive Audit Protocol)
8. Appendix AI (Term Discrimination and Semantic Stability Protocol)
9. Appendix AJ (Public Release Hardening Addendum)
10. Appendix AK (Computability vs Realizability Bridge)
11. Appendix AL (Source-Coupling Integrity and Generative-Source Traceability)
12. Appendix AM (Physical/Causal Admissibility Evidence Profile)
13. Appendix AN (Methodological Falsifiability and Dependency Integrity Standard)
14. Appendix AO (Physical Execution Boundary)
15. Appendix T (Qualification Continuity Check)
16. Appendix U (Outcome Requalification and Evidence Basis)
17. Appendix AE (Layer-Discipline, Falsification, and Adaptive Reserve Addendum)
18. Appendix AF (Decision Verdict Semantics, Projection Freeze, Closure Status, Falsification Classes, and Authority-Selection Separation)
19. Appendix AX (Normative and Claim-Authority Clarification Patch): normative within its declared interpretation and claim-authority scope; subordinate to the substantive gate rules, equations and higher-precedence authorities above.

20. Appendix AQ-1 (Configuration-Bound Assurance and Reader Navigation): its configuration-binding, frozen-evaluation, requalification and Open Assurance Dependency requirements are normative clarifications within their declared scope and subordinate to the authorities above. Its abstract notes, process model and reader navigation are informative.
21. Appendix A (Symbols, Domains, and Notation Canon)
22. Appendix E (UCI Interface and Tier-3 Construction Guidance)
23. Appendix G (SGP Integration Binding Interface)
24. Appendix R (Tier 1-3 Reference Test Vectors)
25. Appendix V (Measurement Maturity and Deployment Context Hardening)
26. Appendix W (Schema Ownership, Registries, and Machine-Verifiable Surfaces)
27. Appendix AQ (Current Canonical Machine-State and Maturity Contract): normative for emitted machine-state vocabulary and maturity typing within its declared interface; subordinate to the substantive decision and equation authorities above. Its external-registry availability limitations remain controlling for interchange claims.
28. Appendix X (Reference Calculator, Machine Vectors, and ProofPack-Prep Roadmap)
29. Appendix Y (Precision Patch: Legitimacy, Evidence Status, Validation, and Adoption Profiles)
30. Appendix Z (Integration-Hardening Addendum)
31. Appendix AA (Package-Content Boundary and Tooling-Line Separation)
32. Appendix AB (Reference Structure Protocol)
33. Appendix N (Stewardship Pack Canon)
34. Appendix F (Failure Modes and Anti-Gaming Controls)
35. Appendix J (Pre-Registration Templates and Scoring Definitions)
36. Appendix AD (49-Cell Welfare Dictionary): canonical interpretive reference for cell meanings, scoring language, and reviewer checks. Subordinate to Appendix B equations and to all normative gates; governs scoring interpretation and reviewer literacy only. It does not control admissibility, thresholds, or claim boundaries.

Informative appendix status. Appendices P, P-2, Q, and S are worked-example, failure-path, sensitivity, or explanatory surfaces. Appendix AC houses this Specification Contract. These appendices may demonstrate and clarify current rules but do not override the main text or the controlling gate equations in Appendices B-D.

37. Examples, templates, and illustrative values (non-normative; includes Appendix K unless a future version explicitly elevates a named subsection)
38. Appendix O (Conformance Reference Map; informative pointer layer only)

Note (Informative). Version history and release notes are recorded in Appendix M and do not control normative meaning.

Priority-direction note (Normative clarification). Item 1 is the highest controlling authority in this list. Higher-numbered items are subordinate to lower-numbered items. Appendix O is informative and cannot override any normative rule.

Release-authority scope (Normative). Appendix RELEASE controls package identity, component editions, build identity, artifact-hash references, source ownership and provenance. It does not change substantive

decision semantics or promote an unavailable external artifact into verified conformance. Identity metadata does not override the semantic precedence order above.

0.1A v13.0 Current Authority and Release-Identity Rule (Normative)

RippleLogic v13.0 is governed first by main-text rules marked MUST, SHALL, or PROHIBITED; then by Appendix B equations, Appendix C rights rules, Appendix D TRC rules, Appendix H PCC and audit-flag rules, and the remaining controlling appendices in the order declared by Section 0.1. Companion standards control only their bounded interfaces and remain subordinate to the Canon on conflict.

The delivered release identity is the exact component matrix in Section 0.2 together with the supplied build manifest and SHA-256 ledger. Repository VERSION_MANIFEST.yaml and release/release_manifest.yml require separate synchronization before promotion. Any identity stated in title pages, headers, footers or component maps must agree with Appendix RELEASE; a mismatch invalidates a release-integrity claim. No publication or repository conformance is inferred.

Historical or lineage references remain valid only when explicitly marked HISTORICAL, LINEAGE, SUPERSEDED, or PRIOR RELEASE. Unmarked version references are current-release claims and MUST match the active matrix.

Parenthetical appendix origin stamps governed by Appendix M.5A are permitted historical lineage markers; they do not change current edition or source authority.

Authority surface	v13.0 role
Canon main text and controlling appendices	Normative semantics, equations, gates, selection, records and claim boundary
SGP v8.8	MPS/FPP/GPR/SPR/ICP/RMCP semantics only; no cascade override
Named companion standards	Bounded operational interfaces; Canon controls on conflict
Machine registries and schemas	Exact tokens and record validation; cannot change Canon meaning
Workbooks, examples, Primer and Public Introduction	Informative or worked-run surfaces; cannot create authority or stronger claims

0.2 Tier contract and claims (v13.0 Canon line; Core Release 2026.09 package line) Tier 1-3: current procedural specification within the declared profile, tested scope and unresolved-evidence boundary. Core operating-document pins are Canon v13.0 / SGP v8.8 / ripple.md v5.8 / Agent System v13.0 / Primer v4.7 / Aligners Sheet v5.9 / Cascade Standard v2.9 / CSV Gate Standard v2.7 / Reproducibility Standard v1.7 / WDBIP v1.9 / RLS Validation Protocol v2.9 / Public Introduction v13.0 / PC-AEP v2.6 / MFDI v2.6 / Source-Coupling Integrity Standard v2.6. The supplied build manifest identifies this Core 15 package. RLS Validation Workbook v0.4 and repository machine interfaces, schemas, validators and release registries are external dependencies, not newly supplied or verified artifacts. Any conformance claim MUST pin the exact release ID, artifact versions, release manifest, and SHA-256 hashes. Tier 4: PROHIBITED to claim in this release (design target only). Any Tier-4 content remains informational architecture describing what becomes claimable only after ProofPack publication and independent replayability.

0.3 Conformance vocabulary (Normative) RippleLogic uses distinct predicates; implementations MUST NOT collapse them.

- Rights-admissible: passes NCRC.
- Admissible: the option has `RG_status` in `{RG_SUPPORTED, RG_NARROWED}`, passes NCRC, and has `TRC_status` in `{TRC_PASS, TRC_NOT_TRIGGERED}`. Formally: $\text{Admissible}(a) := \text{RG_qualified}(a) \text{ AND NCRC}(a) \text{ AND TRC_qualified}(a)$. This is claim-bounded admissibility, not execution authority.
- Selectable: admissible and has `CSV_status` in `{CSV_PASS, CSV_PASS_WITH_CONTROLS, CSV_NOT_MATERIAL}`. Controls are part of the option; removing them reopens CSV.
- Selected: the final chosen option from the selectable set, either by `ALLOW_FRAMEWORK_SELECTION` when the framework selection path is decisive or by a separately disclosed `AuthoritySelectionRecord` when the framework refuses deterministic selection. A selected option MUST NOT be represented as a unique framework verdict when `FrameworkVerdict` is `REFUSE_DETERMINISTIC_SELECTION`.

0.4 Consolidated symbols and core objects (Normative; v13.0, symbol set extended for the current edition) Terminology note (Normative). Since v11.0 and carried forward in v12.0, `u` indexes Union Scopes (short: Scopes). “Union(s)” remains an allowed legacy narrative alias. This clarification SHALL NOT change the meaning of `u` or cascade semantics. Machine-readable surfaces SHOULD prefer `UnionScope` / `UnionScopes` as the canonical identifier family; `Union` remains a permitted human-readable alias in prose and UI text.

Notation scope rule (Normative clarification). A symbol must be interpreted with its indices, field path and owning equation. In Appendix D.7, $\beta = 1 - \alpha$ is tail mass and δ is boundary probability mass, not the Section 5 saturation β or Section 10 discrimination δ . $q(u,d)$, q_u and q_k respectively denote effective cell weight, declared scope prominence and local confidence-contribution weight; Q , Q_{cell} and an example-specific Q_u are distinct objects. Likewise spectral

rho(K), scoped rights tolerances and residual shares are not interchangeable; h_k is welfare inclusion, not h_sev. Alpha is also scoped: the unsubscripted alpha in TRC is the tail level, alpha_H/alpha_F/alpha_R/alpha_E are UCI component weights, and alpha_u in a conserved-allocation record is an allocation share. None is an alias for the others. Implementations must preserve these qualified names and must not infer a shared parameter from a reused glyph.

Supplemental symbols: z is raw system state; e_effect is an identified consequence object; I is normalized signed impact; C_ud, C_u and C_d are final weighted-cell, scope and dimension contributions; Q_cell is local confidence-aggregation mass, not the global RLS mass Q. d_prime is a source-dimension index. chi is a structural susceptibility parameter, not a protected-subgroup index. M_dyn and B_dyn are state-dynamics and control-input matrices; xi is disturbance; Id is the identity matrix. T_control_upper, T_harm_lower and margin are the governed timing bounds in §9.9A. eta and psi, where used in a control model, denote effectiveness and residual transmissivity under that model, not new welfare weights. Kernel bounds k_max and k_use_min retain the numerical values formerly displayed as kappa_max and kappa_use_min; existing machine-field names remain unchanged.

Table 0-1. Indices and core symbols

Symbol	Meaning
$u \in \{1, \dots, 7\}$	Union Scope (Scope) index (legacy alias: union)
$d \in \{1, \dots, 7\}$	Dimension index
$a \in O$	Candidate option
$s \in S$	Scenario index
$r \in R$	Right index
k	Impact instance index
$g \in G_{\{u,d\}}$	Protected subgroup index (for rights checking)

Table 0-2. Core sets

Set	Definition
$U = \{U_1, \dots, U_7\}$	Operational Scopes
$D = \{D_1, \dots, D_7\}$	Welfare dimensions
R	Rights set: {LIFE, BODY, LBTY, NEED, DIGN, PROC, INFO, ECOL}
$C_r \subseteq U \times D$	Rights coverage sets
$C_{cat} \subseteq U \times D$	Catastrophe cell set for TRC

Table 0-3. Welfare impact objects (all options a)

Symbol	Meaning	Range
$I_{dir}(u,d,a)$	Direct impact (post-saturation)	[-1,+1]
$I_{prop}(u,d,a)$	Propagated impact (post-saturation)	[-1,+1]
$I_{rights}(u,d,a)$	Rights worst-off subgroup impact	[-1,+1]

Table 0-4. Kernel, weights, and tail-risk symbols

Symbol	Meaning	Constraint / Notes
$K \in \mathbb{R}^{49 \times 49}$	Ripple kernel matrix (sparse)	Convention: K_{ij} maps source j to target i (target-row, source-column)
w_u	Union weights	$w_u \geq 0, \sum_u w_u = 1$
v_d	Dimension weights	$v_d \geq 0, \sum_d v_d = 1$
$m(u,d)$	Applicability mask	$m \in \{0,1\}$ (RLS aggregation only)
p_s	Scenario probabilities	$p_s \geq 0, \sum_s p_s = 1$
ω_c	Catastrophe weights	$\omega_c \geq 0, \sum_{c \in C_{cat}} \omega_c = 1$
α	CVaR tail level	$\alpha \in (0,1)$
τ_{TRC}	Corridor threshold	$\tau_{TRC} \in (0,1]$
$\kappa(u,d)$	Governed RLS cell interaction multiplier	$0.5 \leq \kappa \leq 1.5$; §10.1. Not k_{raw} or k_{shrunk} in Appendix K.3.
$q(u,d)$	Effective comparison-cell weight	$q = w_u v_d m \kappa$; nonnegative; §10.1.
Q	Active effective comparison mass	$Q = \sum q$; $Q > 0$ required for RLS; not Q_{cell} or an example Q_u .
$\sigma_{RLS(a)}$	Declared option-score uncertainty quantity	Nonnegative; interpretation and dependence treatment under §10.3; not automatically a statistical standard deviation.

Table 0-5. Structural metrics

Symbol	Meaning	Range
UCI_u	Union Coherence Index for scope u	$[0,1]$
$\Delta UCI_u(a)$	$UCI_u(a) - UCI_u(\text{baseline})$	$[-1,+1]$
HOI	Hollowing-Out Index (monitoring diagnostic)	unbounded

0.5 Version commitments (Normative; v13.0) The canonical welfare impact scale is [-1,+1] with Baseline-Zero Rule (Section 5). Optional UI scale conversion (non-normative): If a user interface uses a percent-like scale, map the canonical cell impact $I(u,d,a) \in [-1,+1]$ to UI points via: $\text{points} = 100 \times I$. All admissibility gates (NCRC, TRC, CSV) and RLS calculations **MUST** be performed on the canonical [-1,+1] scale, or be converted back exactly before computation. Tier 1-3 TRC uses bounded-impact loss (Section 8). Tier 4 is design target only; no Tier-4 claims permitted until ProofPack is publicly replayable. UCI components that are genuinely inapplicable are marked NA, excluded, and the remaining weights are renormalized. For Self (U_1), Equity is ordinarily NA unless a governed within-self equity instrument is declared. For Biosphere (U_7), Equity is ordinarily NA unless a governed ecological-distribution or ecological-justice instrument is declared. Missing evidence for an applicable component is UNKNOWN, not NA or a fixed perfect score.

Patch-line interpretation (Normative clarification). References in section headings or appendix titles to v12.0, v12.1, v12.2, v12.4 are origin and lineage labels when explicitly marked historical, lineage, carried-forward, or superseded. Unless a rule explicitly names a different controlling version, active normative rules in this artifact apply under the v13.0 line. Same-version hardening **MUST** preserve the cascade architecture unless an explicit versioned change states otherwise.

0.6 Single Canonical Artifact Rule (Normative; v13.0)

This is the separately identified RippleLogic Canon v13.0 in the MathGov/RippleLogic v13.0 package. Identity requires the current filename and manifest hash. This local edition does not silently replace historical runs or imply that an external repository has been published or synchronized.

Baseline-preservation rule. Preserve the recovered v12.8-rc2 package, the audited v12.7 lineage and all prior PCCs unchanged. The current amendments apply only to separately pinned v13.0 files and the named current tests. A legacy run remains a legacy run unless deliberately migrated and rerun; modified bytes require a new identity and hash.

Integrity note (Informative). The canonical hash **MUST** be published in a sibling release manifest rather than embedded in the document body, because any edit changes the byte stream.

Canonical source designation. For this delivered package the Core_15 DOCX masters control edited prose, equations and tables, subject to the Canon/SGP source hierarchy. Sources contains mechanically generated reading projections and text maps; Reading_PDFs contains rendered mirrors. The prior repository Markdown source remains external and unchanged. Any later repository publication requires reviewed integration, parity checks and its own manifest. Appendix RELEASE states current local source precedence and the scope of supplied reference code and schemas.

0.7 License, Attribution, and Marks (Informative)

License (Informative). The v13.0 release package, including this Canon, documentation, code, schemas, templates, and test fixtures, is distributed under the Apache License 2.0 unless a future file carries an explicit SPDX exception. Attribution: cite this document by canonical filename + SHA-256 and credit James McGaughan (originator / system architect). Marks: "RippleLogic" and "MathGov" may be used to describe compatible implementations, but trademark or branding use is governed separately by the project's published brand guidelines when available.

0.8 Release Packaging and Integrity Manifests (Informative)

Recommended release bundle (Informative). Publish (i) the canonical DOCX/PDF, (ii) a SHA-256 manifest file (for example, RippleLogic_<version>.sha256) listing hashes for each artifact, (iii) a machine-readable release manifest when available (for example YAML), and (iv) optional signatures. The manifest is the stable anchor for byte identity, while the document remains human-readable and editable without self-referential hashing issues. To avoid self-reference problems, a checksum manifest SHOULD hash the release artifacts and MAY omit its own file unless a separate detached signature or higher-level release record governs it.

APPENDIX AD: 49-CELL WELFARE DICTIONARY (CANONICAL INTERPRETIVE REFERENCE)

Tooling warning (Normative for implementation). Tooling MUST source rights coverage from Section 7.2 / Appendix C.2 and catastrophe cells from Section 8.2, never from Appendix AD icons. Appendix AD icons are interpretive cues, not machine-authoritative gate definitions.

Status and boundaries. Appendix AD is the canonical interpretive reference for welfare-cell meanings, baseline-relative scoring language, evidence literacy, and reviewer checks. It is binding for interpreting what each Scope × Dimension cell is about. Its examples, indicator lists, and gate-relevance tags are illustrative and non-exhaustive. Appendix AD does not modify the cascade, equations, NCRC, TRC, CSV, RLS, UCI/HOI, HDW, SGP, PCC, RSP/RefStructRecord, thresholds, or claim boundaries. Gate tags are contextual cues for evaluators, not automated verdicts; admissibility is determined only by running the full cascade against declared evidence in the governed system.

The welfare matrix is the surface RLS ranks on: seven union scopes (U1–U7) × seven welfare dimensions (D1–D7), giving forty-nine accountability domains.

The 49-cell welfare matrix is a proposed minimal conceptual covering set for systematic welfare-impact review. It is not yet an empirically validated universal measurement ontology, and the seven dimensions are not assumed to be statistically independent. Each cell is scored on a canonical –1 to +1 scale, relative to a declared baseline. A cell score is a claim, not a vibe: it asserts what changes, for whom, how strongly, on what evidence, and under what uncertainty.

Scale. +1 strong benefit, +0.5 moderate benefit, 0 no material change from baseline — never “unknown”, –0.5 moderate harm, –1 severe harm. Where an impact is not known, the cell is marked unknown, disclosed, and escalated where gate-relevant.

Ordering. Welfare scoring is the last step, not the first. The cascade is sequenced as Reality Grounding → Rights Floor → TRC → CSV → RLS — and an option is ranked across these cells only after it has cleared the rights, tail-risk, and CSV gates. No welfare score rescues an option that failed an earlier gate.

Gate-relevance key. Tags below indicate where a cell may become relevant to a gate in a real run; they are contextual cues, not fixed classifications, and do not replace a canonical evaluation. Canonical rights coverage set C_r and declared catastrophe cell set C_cat, not the icon labels alone, control non-maskability and gate status.

 Rights-relevance cue — may require NCRC review when the cell lies in the canonical or governed rights coverage set C_r, or when the PCC declares a rights-relevant extension. The canonical C_r, not this icon alone, controls rights-covered status and non-maskability.

 Tail-risk relevance cue — may require TRC review when the cell lies in the declared catastrophe cell set C_cat, or when catastrophe relevance is otherwise triggered. The declared C_cat and TRC trigger rules, not this icon alone, control catastrophe-cell status and non-maskability.

 Containment-relevant — a local gain may degrade the wider scopes the decision depends on; checked before RLS.

- Ordinary welfare — ranked by RLS once RF/NCRC, TRC, and CSV are cleared.
- ? Uncertain — evidence weak or contested; mark, disclose, and escalate if gate-relevant.

AD.1 Tabular dictionary

Appendix AD now presents the complete 49-cell dictionary in two layers: first a full 7 x 7 orientation map, then detailed Union Scope tables for meaning, positive/negative movement, evidence, gate cues, examples, and reviewer checks. This preserves the full Appendix AD cell structure while making the text easier to read in DOCX, PDF, and GitHub views.

Reading rule. A cell score is baseline-relative. + Benefit and - Harm describe the direction of movement. Evidence and gate cues identifies typical evidence and possible gate relevance. Example and reviewer check gives one example and the minimum challenge a reviewer should make. Gate cues are not verdicts; the canonical cascade remains controlling.

Machine-readable companion. The Canon keeps Appendix AD in human-readable split tables for document clarity. A synchronized one-row-per-cell companion is also shipped at docs/canon/AD_49_Cell_Welfare_Dictionary.csv. The CSV preserves the older wide-record strength for tooling, audits, and implementation work by keeping each welfare cell on one row with separate fields for meaning, benefit movement, harm movement, evidence, gate cues, example, and reviewer check. The Canon text remains controlling if any companion representation diverges.

AD.1A Full 49-cell map (all coordinates preserved)

This map restores the full 7 x 7 Appendix AD view in one table. Each of the 49 cells appears once. Use this table for orientation; use the detailed tables that follow for meanings, positive/negative movement, evidence, gate cues, examples, and reviewer checks.

Union Scope	D1 Material	D2 Health	D3 Social	D4 Knowledge	D5 Agency	D6 Meaning	D7 Environment
U1 Self	U1/D1 - Personal Resources	U1/D2 - Personal Health	U1/D3 - Relationships	U1/D4 - Learning	U1/D5 - Personal Agency	U1/D6 - Purpose	U1/D7 - Local Conditions
U2 Household	U2/D1 - Household Resources	U2/D2 - Household Health	U2/D3 - Family Cohesion	U2/D4 - Shared Learning	U2/D5 - Household Agency	U2/D6 - Family Meaning	U2/D7 - Home Environment
U3 Community	U3/D1 - Local Resources	U3/D2 - Community Health	U3/D3 - Community Cohesion	U3/D4 - Local Knowledge	U3/D5 - Community Agency	U3/D6 - Shared Meaning	U3/D7 - Local Ecology
U4 Organization	U4/D1 - Organizational Resources	U4/D2 - Organizational Safety	U4/D3 - Organizational Culture	U4/D4 - Knowledge Systems	U4/D5 - Role Agency	U4/D6 - Mission Coherence	U4/D7 - Operational Footprint
U5 Polity	U5/D1 - Public Resources	U5/D2 - Population Health	U5/D3 - Legitimacy	U5/D4 - Information Access	U5/D5 - Civil Agency	U5/D6 - Public Meaning	U5/D7 - Public Environment
U6 Humanity / Global Coordination	U6/D1 - Global Resources	U6/D2 - Human Safety	U6/D3 - Cooperation	U6/D4 - Civilizational Knowledge	U6/D5 - Collective Agency	U6/D6 - Shared Purpose	U6/D7 - Planetary Conditions
U7 Biosphere	U7/D1 - Life-Support Resources	U7/D2 - Ecosystem Health	U7/D3 - Biotic Relations	U7/D4 - Ecological Knowledge	U7/D5 - Resilience Capacity	U7/D6 - Life Continuity	U7/D7 - Ecological Integrity

AD.1B Detailed 49-cell dictionary

U1 Self - meaning and movement

Cell	What this cell measures	+ Benefit	- Harm
U1/D1 - Personal Resources (label: personal resources)	An individual's economic security: income, savings, housing stability, and access to the material means of a decent life.	More stable income, affordable housing, reduced debt, secured access to essentials.	Job loss, eviction risk, debt spirals, loss of essential goods or services.
U1/D2 - Personal Health (label: personal health)	The individual's physical and mental health, safety, and bodily integrity.	Lower injury/illness risk, better access to care, improved mental health, safer conditions.	Injury, illness, untreated conditions, exposure to hazards, deteriorating mental health.
U1/D3 - Relationships (label: relationships)	The quality and stability of a person's close personal relationships and social connection.	Stronger bonds, reduced isolation, more supportive networks.	Isolation, relationship breakdown, forced separation, loss of support.
U1/D4 - Learning (label: learning)	The individual's access to education, skills, and the capacity to learn and develop.	Better access to learning, skill gains, expanded opportunity.	Lost access to education, deskilling, barriers to learning.
U1/D5 - Personal Agency (label: autonomy)	The person's real ability to choose, refuse, consent, act, and shape their own life path.	More autonomy, better options, stronger consent, freedom from coercion and surveillance.	Coercion, manipulation, surveillance, forced dependency, loss of meaningful choice.
U1/D6 - Purpose (label: purpose)	The individual's sense of meaning, dignity, and worth in their life and work.	Greater sense of purpose, recognition, dignity, and self-respect.	Humiliation, alienation, loss of role or recognition, indignity.
U1/D7 - Local Conditions (label: local conditions)	The immediate physical environment a person lives in: air, noise, water, hazards, and surroundings.	Cleaner air/water, less noise, fewer local hazards, healthier surroundings.	Pollution exposure, noise, contamination, environmental hazards near the person.

U1 Self - evidence and review

Cell	Evidence and gate cues	Example and reviewer check
U1/D1 - Personal Resources (label: personal resources)	Evidence: income and employment records; housing stability / eviction data; debt and arrears; cost-of-living and benefit uptake; Gate classes: Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Ordinary welfare in most cases; can escalate to rights-covered where a subsistence or social-minimum floor is legally recognised.	Example: A wage policy that lifts low-income workers above a local living-wage line improves U1/D1.; Reviewer check: A reviewer checks whether the income/housing change is measured against a declared baseline cohort, not anecdote.
U1/D2 - Personal Health (label: personal health)	Evidence: morbidity / mortality data; care access and waiting times; injury and exposure records; validated mental-health measures; Gate classes: Rights Floor; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Frequently rights-covered: bodily integrity and access-to-care floors can trigger NCRC. Severe individual harm escalates regardless of aggregate benefit.	Example: A treatment-access change that removes a barrier for a chronic-illness group improves U1/D2.; Reviewer check: A reviewer asks whether a claimed health gain rests on outcome data or only on inputs (e.g. clinics built ≠ health improved).
U1/D3 - Relationships (label: relationships)	Evidence: loneliness / connection surveys; support-network mapping; separation or displacement records; Gate classes: Ordinary welfare. Uncertainty: Explicitly elevated. Interpretation: Usually ordinary welfare; measurement is often soft, so uncertainty should be marked explicitly.	Example: A relocation program that keeps support networks intact protects U1/D3.; Reviewer check: A reviewer challenges whether a relationship claim is evidenced or inferred from a proxy like attendance.
U1/D4 - Learning (label: learning)	Evidence: enrolment and completion; skills assessments; access and dropout data; Gate classes: Ordinary welfare. Uncertainty:	Example: A reskilling program for displaced workers improves U1/D4.;

Cell	Evidence and gate cues	Example and reviewer check
	Standard assessment. Interpretation: Ordinary welfare; can touch rights where a basic-education entitlement exists.	Reviewer check: A reviewer asks whether learning outcomes are demonstrated or only opportunity offered.
U1/D5 - Personal Agency (label: autonomy)	Evidence: consent quality; availability of alternatives / opt-out; coercion complaints; mobility and accessibility data; applicable legal protections; Gate classes: Rights Floor; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Often rights-covered. Triggers NCRC when liberty, bodily integrity, due process, dignity, or information rights are impaired — non-compensatory regardless of welfare gains elsewhere.	Example: A transport policy that gives disabled residents reliable independent mobility improves U1/D5.; Reviewer check: A reviewer tests whether consent was genuine and revocable, or merely formal.
U1/D6 - Purpose (label: purpose)	Evidence: dignity / recognition measures; role and identity surveys; qualitative testimony; Gate classes: Rights-relevance cue; Ordinary welfare. Uncertainty: Explicitly elevated. Interpretation: Rights review is required when the underlying effect instantiates dignity, discrimination, coercion, conscience, or another protected interest routed through Appendix C.2 or a governed extension; otherwise this remains ordinary D6 welfare. Often hard to measure — mark uncertainty.	Example: A program that treats welfare recipients with dignified, non-stigmatising processes supports U1/D6.; Reviewer check: A reviewer distinguishes a genuine dignity effect from sentiment unsupported by experience data.
U1/D7 - Local Conditions (label: local conditions)	Evidence: local air/water/noise monitoring; hazard proximity; exposure records; Gate classes: Rights-relevance cue; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Ordinary D7 welfare unless the underlying exposure is separately routed to Health, Basic Needs/habitability, Dignity, Agency, or another protected interest under Appendix C.2 or a governed extension.	Example: A factory buffer-zone rule that cuts a household's pollutant exposure improves U1/D7.; Reviewer check: A reviewer checks measured exposure change against the affected location, not a city-wide average.

U2 Household - meaning and movement

Cell	What this cell measures	+ Benefit	- Harm
U2/D1 - Household Resources (label: household resources)	The economic security of the household unit: shared income, assets, housing, and ability to meet needs.	Improved household finances, housing security, reduced shared debt.	Household impoverishment, housing loss, debt, inability to meet needs.
U2/D2 - Household Health (label: household health)	The collective physical and mental health and safety of household members, including dependents.	Better shared health, safer home, improved care for dependents and caregivers.	Household illness, unsafe housing, caregiver burnout, harm to dependents.
U2/D3 - Family Cohesion (label: family cohesion)	The strength, stability, and supportive quality of relationships within the household.	Stronger family bonds, stability, reduced conflict, kept-together units.	Family breakdown, forced separation, household conflict, instability.
U2/D4 - Shared Learning (label: shared learning)	The household's collective access to information, education, and learning capacity.	Better shared access to education and information, intergenerational learning.	Loss of household educational access, information poverty.
U2/D5 - Household Agency (label: household agency)	The household's collective ability to make decisions, plan, and control its own circumstances.	Greater household self-determination, planning ability, freedom from external control.	Loss of household control, externally imposed decisions, dependency.
U2/D6 - Family Meaning (label: family meaning)	The household's shared sense of identity, belonging, dignity, and purpose.	Stronger shared identity, belonging, dignity, and continuity.	Loss of household identity, stigma, indignity, eroded belonging.

Cell	What this cell measures	+ Benefit	- Harm
U2/D7 - Home Environment (label: home environment)	The physical environmental quality of the home and its immediate setting.	Healthier home conditions, reduced indoor hazards, better surroundings.	Indoor pollution, mould, hazards, degraded immediate environment.

U2 Household - evidence and review

Cell	Evidence and gate cues	Example and reviewer check
U2/D1 - Household Resources (label: household resources)	Evidence: household income and assets; housing tenure; arrears and food-security data; Gate classes: Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Ordinary welfare; escalates toward rights where dependents (children, disabled members) face a subsistence floor.	Example: A childcare subsidy that frees household income improves U2/D1.; Reviewer check: A reviewer asks whether the household, not just the earner, is the measured unit.
U2/D2 - Household Health (label: household health)	Evidence: household health surveys; home-safety inspections; caregiver-strain measures; dependent outcome data; Gate classes: Rights Floor; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Frequently rights-covered where children or dependents are affected; child welfare can trigger NCRC. Caregiver harm is easily hidden — surface it.	Example: A home-care support program that reduces caregiver strain and improves a dependent's health raises U2/D2.; Reviewer check: A reviewer checks that dependents and caregivers are scored separately, not blended into one household figure.
U2/D3 - Family Cohesion (label: family cohesion)	Evidence: family-stability surveys; separation / removal records; conflict indicators; Gate classes: Rights Floor; Ordinary welfare. Uncertainty: Explicitly elevated. Interpretation: Can be rights-covered (family unity, child rights); otherwise ordinary. Measurement is soft — mark uncertainty.	Example: A policy that avoids unnecessary family separation in housing transitions protects U2/D3.; Reviewer check: A reviewer asks whether a cohesion claim rests on outcomes or on assumed effects of a program.
U2/D4 - Shared Learning (label: shared learning)	Evidence: household connectivity / device access; children's schooling continuity; adult-learning access; Gate classes: Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Ordinary welfare; touches rights where basic-education entitlements for children apply.	Example: A home-broadband program that enables children's remote schooling improves U2/D4.; Reviewer check: A reviewer checks access actually reaches the household, not just the area.
U2/D5 - Household Agency (label: household agency)	Evidence: decision-autonomy surveys; dependency / conditionality of support; forced-relocation data; Gate classes: Rights Floor; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Can be rights-covered (family privacy, autonomy); otherwise ordinary.	Example: A flexible-benefit design that lets households allocate support to their own priorities raises U2/D5.; Reviewer check: A reviewer tests whether household choice is real or constrained by conditions.
U2/D6 - Family Meaning (label: family meaning)	Evidence: belonging / identity surveys; stigma indicators; qualitative testimony; Gate classes: Ordinary welfare. Uncertainty: Explicitly elevated. Interpretation: Usually ordinary welfare; can touch dignity rights. Soft to measure — mark uncertainty.	Example: A culturally-respectful resettlement process preserves U2/D6.; Reviewer check: A reviewer separates a real meaning effect from program rhetoric.
U2/D7 - Home Environment (label: home environment)	Evidence: housing-quality inspections; indoor air / damp measures; hazard records; Gate classes: Rights-relevance cue; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Ordinary D7 welfare; habitability, health, dependent-protection, or basic-needs effects must be separately routed to the applicable rights-covered dimension.	Example: A retrofit program that removes household mould improves U2/D7.; Reviewer check: A reviewer checks measured home conditions, not self-report alone.

U3 Community - meaning and movement

Cell	What this cell measures	+ Benefit	- Harm
U3/D1 - Local Resources (label: local resources)	The shared economic resources, infrastructure, and services available to a local community.	Better local infrastructure, services, and economic opportunity.	Loss of services, infrastructure decay, local economic decline.
U3/D2 - Community Health (label: public health)	The collective physical and mental health and safety of the local population.	Lower community disease/injury burden, better local care access, safer environment.	Local disease spread, unsafe conditions, reduced care access, environmental health harm.
U3/D3 - Community Cohesion (label: trust and cohesion)	The quality of trust, cooperation, belonging, safety, and relational fabric within a community.	Increased trust, reduced isolation, better cooperation, stronger neighbourhood resilience.	Polarisation, exclusion, displacement, fear, social fragmentation.
U3/D4 - Local Knowledge (label: local knowledge)	The community's shared knowledge, information access, and local expertise.	Better local information access, preserved local/Indigenous knowledge, stronger civic literacy.	Information loss, erosion of local knowledge, misinformation spread.
U3/D5 - Community Agency (label: community agency)	The community's collective capacity to organise, participate, and influence decisions affecting it.	Stronger local voice, participation, self-organisation, and contestability.	Disenfranchisement, suppressed organising, exclusion from decisions.
U3/D6 - Shared Meaning (label: shared meaning)	The community's shared identity, culture, dignity, and sense of collective purpose.	Stronger shared identity, cultural continuity, pride, and belonging.	Cultural erosion, stigma, loss of identity, collective humiliation.
U3/D7 - Local Ecology (label: local ecology)	The ecological condition of the community's local environment: green space, biodiversity, water, land.	Restored habitat, cleaner local environment, protected green space.	Habitat loss, local pollution, degraded green space, contamination.

U3 Community - evidence and review

Cell	Evidence and gate cues	Example and reviewer check
U3/D1 - Local Resources (label: local resources)	Evidence: local service provision; infrastructure condition; local employment / business data; Gate classes: Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Ordinary welfare; containment-relevant when local gains are extracted at the cost of wider scopes, or when shared infrastructure is hollowed out.	Example: A community facility upgrade that adds usable shared services improves U3/D1.; Reviewer check: A reviewer asks whether the benefit is genuinely shared or captured by a subgroup.
U3/D2 - Community Health (label: public health)	Evidence: local public-health surveillance; care-access data; injury / safety statistics; environmental-health monitoring; Gate classes: Rights Floor; Tail Risk; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Ordinary welfare normally; rights-relevant under disproportionate harm or care-access floors; tail-relevant where local outbreaks or contamination could cascade.	Example: A clean-water intervention that cuts a community's disease burden improves U3/D2.; Reviewer check: A reviewer checks whether subgroup harms are masked by a favourable community average.
U3/D3 - Community Cohesion (label: trust and cohesion)	Evidence: community trust surveys; participation rates; crime / fear-of-crime data; displacement records; Gate classes: Rights Floor; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Containment-relevant because eroded social fabric degrades the wider system; rights-relevant if discrimination, exclusion, or targeted domination appears.	Example: A public-space renovation that increases safe cross-group interaction improves U3/D3.; Reviewer check: A reviewer asks whether cohesion gains for some came via exclusion of others.

Cell	Evidence and gate cues	Example and reviewer check
U3/D4 - Local Knowledge (label: local knowledge)	Evidence: information-access measures; local-knowledge preservation records; misinformation prevalence; Gate classes: Containment / Structural Viability; Ordinary welfare. Uncertainty: Explicitly elevated. Interpretation: Ordinary welfare; containment-relevant where an epistemic commons is degraded. Can be rights-relevant for Indigenous knowledge.	Example: A program documenting local ecological knowledge with community consent improves U3/D4.; Reviewer check: A reviewer checks consent and ownership of the knowledge claimed as preserved.
U3/D5 - Community Agency (label: community agency)	Evidence: participation in local governance; consultation quality; association / assembly freedom; Gate classes: Rights Floor; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Often rights-covered (assembly, association, participation); containment-relevant where civic capacity underpins system stability.	Example: A binding community-consultation requirement on local development raises U3/D5.; Reviewer check: A reviewer tests whether participation was substantive or a procedural formality.
U3/D6 - Shared Meaning (label: shared meaning)	Evidence: cultural-participation data; heritage continuity; community-identity surveys; Gate classes: Rights-relevance cue; Ordinary welfare. Uncertainty: Explicitly elevated. Interpretation: Ordinary D6 welfare; conscience, discrimination, coercion, exclusion, information, or agency effects must be separately routed to the applicable rights-covered dimension or a governed extension. Soft to measure — mark uncertainty.	Example: Protecting a culturally significant site during redevelopment preserves U3/D6.; Reviewer check: A reviewer separates a genuine cultural effect from a symbolic gesture.
U3/D7 - Local Ecology (label: local ecology)	Evidence: local biodiversity / green-space data; pollution monitoring; land-use change; Gate classes: Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Containment-relevant (local ecology supports wider systems); tail-relevant where degradation is irreversible or contamination could cascade.	Example: A watershed-protection rule that prevents local contamination improves U3/D7.; Reviewer check: A reviewer asks whether 'no local effect' was measured or merely assumed.

U4 Organization - meaning and movement

Cell	What this cell measures	+ Benefit	- Harm
U4/D1 - Organizational Resources (label: operating resources)	An organization's operating means: finances, capital, staffing, and capacity to function and deliver.	Stronger finances, capacity, staffing, and operational sustainability.	Insolvency risk, capacity loss, layoffs, operational decline.
U4/D2 - Organizational Safety (label: safety)	The safety of people within and affected by the organization: workers, users, and the public.	Fewer injuries, safer operations, better safety culture and compliance.	Workplace injury, unsafe products/services, public-safety incidents.
U4/D3 - Organizational Culture (label: culture)	The internal relational health of the organization: trust, fairness, inclusion, and cooperation.	Healthier culture, fairer treatment, reduced harassment, stronger cooperation.	Toxic culture, discrimination, harassment, low trust, burnout.
U4/D4 - Knowledge Systems (label: knowledge systems)	The organization's knowledge, data integrity, institutional memory, and learning systems.	Better knowledge retention, data integrity, transparency, and learning.	Knowledge loss, data corruption, opacity, broken institutional memory.
U4/D5 - Role Agency (label: role agency)	The meaningful autonomy, voice, and fair treatment of individuals in their organizational roles.	More voice, fair process, reasonable autonomy, freedom from arbitrary treatment.	Arbitrary power, suppressed voice, exploitation, loss of role autonomy.
U4/D6 - Mission Coherence (label: mission coherence)	The alignment between the organization's stated purpose and its actual conduct and effects.	Stronger purpose-conduct alignment, integrity, and legitimate mission delivery.	Mission drift, purpose-washing, hollow mandate, loss of legitimacy.
U4/D7 - Operational Footprint (label: operational footprint)	The environmental impact of the organization's operations: emissions, waste, resource use,	Lower emissions/waste, reduced resource intensity, smaller harmful footprint.	Pollution, emissions, waste, resource depletion, ecological damage.

Cell	What this cell measures	+ Benefit	- Harm
	land.		

U4 Organization - evidence and review

Cell	Evidence and gate cues	Example and reviewer check
U4/D1 - Organizational Resources (label: operating resources)	Evidence: financial statements; staffing and capacity data; service-delivery metrics; Gate classes: Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Ordinary welfare; containment-relevant when organizational gains are extracted by externalising harm onto community, polity, or biosphere scopes.	Example: A funding change that stabilises a hospital's operating budget improves U4/D1.; Reviewer check: A reviewer checks whether the organizational gain creates uncosted externalities in other cells (the classic 'operating resources up, U7 down' trap).
U4/D2 - Organizational Safety (label: safety)	Evidence: incident and injury rates; safety-compliance audits; near-miss reporting; Gate classes: Rights Floor; Tail Risk; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Rights-covered via labour and safety protections; tail-relevant where failure could cause catastrophic or mass-casualty events.	Example: A maintenance regime that cuts industrial-accident rates improves U4/D2.; Reviewer check: A reviewer checks whether reported safety reflects outcomes or only paperwork compliance.
U4/D3 - Organizational Culture (label: culture)	Evidence: staff-climate surveys; grievance and turnover data; fairness / inclusion measures; Gate classes: Rights Floor; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Ordinary welfare; rights-relevant where discrimination or harassment occurs.	Example: An anti-harassment reform that reduces verified grievances improves U4/D3.; Reviewer check: A reviewer asks whether culture data is independent or self-reported by leadership.
U4/D4 - Knowledge Systems (label: knowledge systems)	Evidence: data-integrity audits; documentation completeness; knowledge-retention metrics; Gate classes: Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Ordinary welfare; containment-relevant where degraded knowledge systems undermine accountability or wider-system reliability.	Example: A records-integrity upgrade that makes decisions auditable improves U4/D4.; Reviewer check: A reviewer tests whether claimed integrity is verifiable or asserted.
U4/D5 - Role Agency (label: role agency)	Evidence: worker-voice mechanisms; grievance outcomes; autonomy / discretion measures; Gate classes: Rights Floor; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Rights-relevant via labour rights and freedom of association; otherwise ordinary.	Example: A grievance process with real remedy improves U4/D5.; Reviewer check: A reviewer checks whether voice mechanisms produce outcomes or are decorative.
U4/D6 - Mission Coherence (label: mission coherence)	Evidence: mission-vs-outcome analysis; stakeholder-trust measures; conduct audits; Gate classes: Containment / Structural Viability; Ordinary welfare. Uncertainty: Explicitly elevated. Interpretation: Ordinary welfare; containment-relevant where mission decay (HOI-adjacent) hollows out the institution. Distinct from the formal HOI diagnostic.	Example: A governance reform that realigns conduct with mandate improves U4/D6.; Reviewer check: A reviewer asks whether coherence is evidenced by conduct or only by restated values.
U4/D7 - Operational Footprint (label: operational footprint)	Evidence: emissions and waste data; resource-use intensity; environmental-compliance records; Gate classes: Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Containment-relevant (footprint degrades wider scopes); tail-relevant where the footprint contributes to irreversible or large-scale harm. The canonical 'local win, wider loss' cell.	Example: A process change that cuts a plant's emissions while sustaining output improves U4/D7.; Reviewer check: A reviewer checks footprint against measured externalities, not offset claims alone.

U5 Polity - meaning and movement

Cell	What this cell measures	+ Benefit	- Harm
U5/D1 - Public Resources (label: public resources)	The polity's fiscal and material capacity: public finances, infrastructure, and provisioning.	Sounder public finances, better infrastructure, sustainable provisioning.	Fiscal crisis, infrastructure decay, depleted public capacity.

Cell	What this cell measures	+ Benefit	- Harm
U5/D2 - Population Health (label: population health)	The health, safety, and mortality outcomes of the whole population within a polity.	Lower population mortality/morbidity, better system-wide care, stronger resilience.	Rising mortality, system overload, reduced care access, public-health crises.
U5/D3 - Legitimacy (label: legitimacy)	The rule of law, due process, accountability, and perceived legitimacy of public institutions.	Stronger rule of law, due process, accountability, and institutional trust.	Erosion of rule of law, impunity, corruption, collapse of institutional trust.
U5/D4 - Information Access (label: information access)	The polity's information environment: access, transparency, free expression, and epistemic integrity.	Better transparency, information access, free expression, healthier epistemic commons.	Censorship, opacity, disinformation, surveillance chilling, epistemic capture.
U5/D5 - Civil Agency (label: civil agency)	People's ability to participate in public life, exercise rights, contest authority, and influence governance.	Stronger voting access, participation, transparency, due process, and contestability.	Suppression, censorship, corruption, procedural exclusion, arbitrary enforcement, reduced civic voice.
U5/D6 - Public Meaning (label: public meaning)	Shared civic identity, dignity, social contract, and non-domination across the polity.	Stronger civic dignity, inclusion, fair recognition, and a credible social contract.	Domination, civic humiliation, exclusionary identity, eroded social contract.
U5/D7 - Public Environment (label: public environment)	The environmental quality and ecological sustainability managed at the polity scale.	Better regional environmental quality, sustainable resource management, reduced systemic pollution.	Regional pollution, resource depletion, ecosystem degradation, environmental injustice.

U5 Polity - evidence and review

Cell	Evidence and gate cues	Example and reviewer check
U5/D1 - Public Resources (label: public resources)	Evidence: public-finance data; infrastructure condition indices; fiscal-sustainability metrics; Gate classes: Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Ordinary welfare; containment-relevant where short-term fiscal wins erode long-term public capacity.	Example: An infrastructure-maintenance program that prevents asset decay improves U5/D1.; Reviewer check: A reviewer checks whether gains are sustainable or borrowed from the future.
U5/D2 - Population Health (label: population health)	Evidence: population mortality / morbidity; health-system capacity; health-equity gaps; outbreak surveillance; Gate classes: Rights Floor; Tail Risk; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Tail-relevant for pandemic / health-system-collapse scenarios (non-averagable); rights-relevant under right-to-health floors; otherwise ordinary welfare.	Example: A vaccination program that lowers population mortality improves U5/D2.; Reviewer check: A reviewer checks whether worst-case scenarios were bounded with a tail measure, not averaged away.
U5/D3 - Legitimacy (label: legitimacy)	Evidence: rule-of-law indices; due-process / judicial-independence measures; corruption indicators; institutional-trust surveys; Gate classes: Rights Floor; Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Strongly containment-relevant: legitimacy is load-bearing for the whole system. Rights-relevant (due process); tail-relevant where collapse risks systemic instability.	Example: A judicial-independence safeguard that resists political capture improves U5/D3.; Reviewer check: A reviewer tests whether legitimacy is measured independently or claimed by the institution being assessed.
U5/D4 - Information Access (label: information access)	Evidence: transparency / FOI metrics; press-freedom indices; disinformation prevalence; surveillance scope; Gate classes: Rights Floor; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Rights-covered (free expression, information rights); containment-relevant because a damaged epistemic commons undermines all other governance.	Example: An open-data and FOI reform that increases verifiable transparency improves U5/D4.; Reviewer check: A reviewer checks whether transparency is real and usable or nominal.
U5/D5 - Civil Agency (label: civil agency)	Evidence: voting access and turnout; procedural-fairness measures; availability of legal remedies; consultation quality;	Example: A reform expanding access to public hearings raises U5/D5; one that

Cell	Evidence and gate cues	Example and reviewer check
civil agency)	censorship reports; Gate classes: Rights Floor; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Frequently rights-covered: tied to due process, liberty, information, dignity, and civic rights — non-compensatory under NCRC. Containment-relevant for system stability.	blocks lawful dissent harms it.; Reviewer check: A reviewer tests whether civic participation is enforceable or merely nominal, and whether any group is selectively excluded.
U5/D6 - Public Meaning (label: public meaning)	Evidence: inclusion / recognition measures; social-contract trust surveys; discrimination indicators; Gate classes: Rights-relevance cue; Containment / Structural Viability; Ordinary welfare. Uncertainty: Explicitly elevated. Interpretation: Dignity, non-domination, discrimination, information, due-process, or agency effects require separate rights routing under Appendix C.2 or a governed extension; the residual public-meaning effect remains ordinary welfare. Containment remains relevant where a fractured social contract destabilises the polity. Soft to measure.	Example: A reconciliation policy that credibly extends recognition to a marginalised group improves U5/D6.; Reviewer check: A reviewer separates symbolic recognition from changes people actually experience.
U5/D7 - Public Environment (label: public environment)	Evidence: regional environmental monitoring; resource-stock data; pollution and emissions; environmental-justice mapping; Gate classes: Rights-relevance cue; Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Tail-relevant where regional harm is irreversible or large-scale; containment-relevant; environmental-justice, health, habitability, dignity, agency, or basic-needs effects require separate routing to the applicable rights-covered dimension or a governed extension.	Example: A regional emissions cap that prevents airshed degradation improves U5/D7.; Reviewer check: A reviewer checks whether environmental-justice distribution was assessed, not just regional totals.

U6 Humanity / Global Coordination - meaning and movement

Cell	What this cell measures	+ Benefit	- Harm
U6/D1 - Global Resources (label: global resources)	Humanity's shared material base: global resource stocks, critical supply chains, and common infrastructure.	More sustainable global resource use, resilient supply chains, protected commons.	Resource depletion, supply-chain fragility, degradation of global commons.
U6/D2 - Human Safety (label: human safety)	The safety and survival of humanity at large, including existential and catastrophic risk.	Reduced catastrophic and existential risk; stronger global safety and resilience.	Increased existential, pandemic, conflict, or technological catastrophic risk.
U6/D3 - Cooperation (label: cooperation)	The capacity of humanity to coordinate, cooperate, and maintain peaceful, stable relations at scale.	Stronger global cooperation, conflict reduction, durable coordination institutions.	Breakdown of cooperation, arms races, institutional collapse, escalating conflict.
U6/D4 - Civilizational Knowledge (label: civilizational knowledge)	Humanity's accumulated knowledge, science, and the integrity and safety of its knowledge systems.	Knowledge preservation, scientific progress, epistemic integrity, responsible knowledge stewardship.	Knowledge loss, scientific capture, epistemic collapse, or proliferation of dangerous knowledge.
U6/D5 - Collective Agency (label: collective agency)	Humanity's capacity for self-determination and legitimate collective decision-making at the species scale.	Stronger legitimate global agency, fair representation, and self-determination.	Concentration of unaccountable power, disenfranchisement of peoples, loss of collective self-determination.
U6/D6 - Shared	Humanity's shared sense of	Stronger shared purpose, long-	Nihilism, short-termism, loss of

Cell	What this cell measures	+ Benefit	- Harm
Purpose (label: shared purpose)	meaning, moral direction, and long-term orientation toward flourishing.	term orientation, intergenerational responsibility.	shared direction, intergenerational betrayal.
U6/D7 - Planetary Conditions (label: planetary conditions)	The condition of planetary systems that support human civilisation and coordinated managing intelligence.	Reduced emissions, improved climate stability, lower systemic ecological risk, stronger global resilience.	Increased climate risk, global pollution, destabilised food/water systems, irreversible planetary degradation.

U6 Humanity / Global Coordination - evidence and review

Cell	Evidence and gate cues	Example and reviewer check
U6/D1 - Global Resources (label: global resources)	Evidence: global resource-stock data; supply-chain resilience metrics; commons-condition indicators; Gate classes: Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Containment-relevant and tail-relevant where depletion or fragility threatens systemic, hard-to-reverse harm at civilisational scale.	Example: A critical-minerals stewardship agreement that reduces depletion risk improves U6/D1.; Reviewer check: A reviewer checks whether 'global' claims rest on global data or extrapolated local figures.
U6/D2 - Human Safety (label: human safety)	Evidence: catastrophic-risk assessments; biosecurity / nuclear-risk indicators; conflict-escalation data; Gate classes: Rights Floor; Tail Risk; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Primarily tail-relevant: catastrophic and existential risk is non-averagable and bounded under TRC, never traded against ordinary welfare. Rights-relevant at the right-to-life floor.	Example: A biosecurity protocol that lowers engineered-pandemic risk improves U6/D2.; Reviewer check: A reviewer checks that tail risk was bounded with a worst-case measure and not diluted by expected-value framing.
U6/D3 - Cooperation (label: cooperation)	Evidence: international-cooperation indices; conflict / treaty-compliance data; coordination-institution health; Gate classes: Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Containment-relevant (cooperation underpins global stability); tail-relevant where breakdown risks large-scale conflict.	Example: A verification regime that sustains an arms-control treaty improves U6/D3.; Reviewer check: A reviewer asks whether cooperation is durable or a fragile short-term arrangement.
U6/D4 - Civilizational Knowledge (label: civilizational knowledge)	Evidence: scientific-integrity measures; knowledge-preservation status; research-governance indicators; Gate classes: Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Containment-relevant; tail-relevant in two directions — catastrophic loss of knowledge, and dangerous diffusion of harmful capability.	Example: An open-but-governed research framework that preserves integrity while limiting dangerous diffusion improves U6/D4.; Reviewer check: A reviewer checks whether knowledge benefits were weighed against misuse and proliferation pathways.
U6/D5 - Collective Agency (label: collective agency)	Evidence: representation / participation in global governance; power-concentration indicators; self-determination measures; Gate classes: Rights Floor; Containment / Structural Viability. Uncertainty: Explicitly elevated. Interpretation: Rights-relevant (self-determination); containment-relevant where illegitimate concentration destabilises the system. Measurement is contested — mark uncertainty.	Example: A governance reform giving affected populations real voice in a global regime improves U6/D5.; Reviewer check: A reviewer tests whether 'collective' agency includes the marginalised or only powerful actors.
U6/D6 - Shared Purpose (label: shared purpose)	Evidence: long-term-orientation indicators; intergenerational-equity measures; qualitative analysis; Gate classes: Containment / Structural Viability; Ordinary welfare. Uncertainty: Explicitly elevated. Interpretation: Mostly ordinary/aspirational; containment-relevant where short-termism erodes long-term stability. Inherently soft — mark uncertainty and avoid overclaiming.	Example: A long-horizon institution (e.g. a future-generations commissioner) can support U6/D6.; Reviewer check: A reviewer flags this cell as especially prone to rhetoric and demands concrete mechanisms.
U6/D7 - Planetary Conditions (label: planetary conditions)	Evidence: emissions data; climate models; planetary-boundary metrics; global food/water-security projections; Gate classes: Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Often catastrophe-relevant: may trigger TRC where irreversible or large-	Example: A high-emissions infrastructure choice may harm U6/D7 even when it benefits a local economy (U3/D1 or U4/D1).; Reviewer check: A reviewer checks whether irreversibility and worst-

Cell	Evidence and gate cues	Example and reviewer check
	scale harm is plausible. Containment-relevant as a system-stability check.	case bounds were modelled, not just expected outcomes.

U7 Biosphere - meaning and movement

Cell	What this cell measures	+ Benefit	- Harm
U7/D1 - Life-Support Resources (label: life-support resources)	The biosphere's foundational resources that sustain life: soil, fresh water, clean air, fertility.	Protected or restored soil, water, and air; sustained natural-resource base.	Soil loss, water depletion, air degradation, exhaustion of life-support systems.
U7/D2 - Ecosystem Health (label: ecosystem health)	The health, function, and viability of ecosystems and the species within them.	Healthier, functioning ecosystems; recovering species and habitats.	Ecosystem dysfunction, species decline, habitat collapse, trophic breakdown.
U7/D3 - Biotic Relations (label: biotic relations)	The integrity of relationships and interdependencies among living systems and between humans and nature.	Restored ecological relationships, sustainable human-nature interactions, protected interdependencies.	Disrupted food webs, broken ecological relationships, harmful human-nature dynamics.
U7/D4 - Ecological Knowledge (label: ecological knowledge)	An explicitly declared epistemic enabling-condition view of biospheric information-bearing conditions. Human understanding or institutional monitoring belongs primarily to the relevant U1-U6 D4 home; a U7 view alone adds no score mass.	Separately evidenced preservation of an ecological information-bearing condition, where distinct from already recorded biodiversity, functioning or monitoring outcomes.	Loss of a separately defined ecological information-bearing condition; lack of human monitoring is instead routed to the appropriate human/institutional D4 record and the control-evidence assessment.
U7/D5 - Resilience Capacity (label: resilience capacity)	The biosphere's capacity to absorb shocks, adapt, and recover from disturbance.	Greater ecological resilience, redundancy, and recovery capacity.	Loss of resilience, reduced redundancy, approach to tipping points and regime shifts.
U7/D6 - Life Continuity (label: life continuity)	The continuity and persistence of life and biodiversity over time, including irreversibility of loss.	Protected biodiversity, prevented extinctions, sustained continuity of life.	Species extinction, irreversible biodiversity loss, broken continuity of lineages.
U7/D7 - Ecological Integrity (label: ecological integrity)	Integrated condition of an ecological life-support system where it is not already scored through the same state's component tokens. The U7 construct rule distinguishes component functioning, relationships, resilience and continuity from an integrated reporting view.	Improved integrated ecological condition with declared evidence and no duplicate component contribution. Component improvements may instead be scored in their primary cells and viewed here without added RLS mass.	Deterioration of the integrated ecological condition, including threshold-relevant damage, with component links retained for rights, TRC and CSV review.

U7 Biosphere - evidence and review

Cell	Evidence and gate cues	Example and reviewer check
U7/D1 - Life-Support Resources (label: life-support resources)	Evidence: soil-health data; freshwater-stock and quality; air-quality baselines; resource-renewal rates; Gate classes: Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Tail-relevant where depletion is irreversible or threatens collapse of life-support systems; containment-relevant. Distinct from	Example: An aquifer-protection rule that prevents irreversible groundwater loss improves U7/D1.; Reviewer check: A reviewer checks whether renewal rates, not just current stocks, were assessed.

Cell	Evidence and gate cues	Example and reviewer check
	U4/D1 'operating resources'.	
U7/D2 - Ecosystem Health (label: ecosystem health)	Evidence: ecosystem-function indicators; species-population trends; habitat-condition assessments; Gate classes: Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Tail-relevant where collapse or extinction is plausible and irreversible; containment-relevant as ecosystems underpin wider scopes.	Example: A fisheries policy that lets a collapsing stock recover improves U7/D2.; Reviewer check: A reviewer checks whether ecosystem-function evidence exists or only a single charismatic indicator.
U7/D3 - Biotic Relations (label: biotic relations)	Evidence: trophic / food-web indicators; interaction-network data; human-impact assessments; Gate classes: Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Containment-relevant: a local win that severs ecological relationships can fail containment even with strong U4/D1 gains; tail-relevant where cascades are plausible.	Example: A factory expansion that damages a watershed may fail containment via U7/D3 even while improving U4/D1.; Reviewer check: A reviewer asks whether relational/cascade effects were modelled or ignored as 'no direct effect'.
U7/D4 - Ecological Knowledge (label: ecological knowledge)	Evidence: the named ecological information-bearing construct, substrate/bearer, baseline, and non-duplication rationale. Human monitoring coverage and early-warning capability remain evidence or primary institutional D4 outcomes, not automatic U7 welfare. Missing monitoring may block a control claim under CSV.	Example: a biodiversity-monitoring programme primarily improves the responsible institution's D4 record. U7 receives a view-only link unless a separate ecological state change is demonstrated. Reviewer check: no duplicate monitoring benefit; absent evidence remains unknown.
U7/D5 - Resilience Capacity (label: resilience capacity)	Evidence: resilience / redundancy indicators; recovery-rate data; threshold-proximity assessments; Gate classes: Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Tail-relevant where loss of resilience brings irreversible regime shifts; containment-relevant as resilience buffers all scopes.	Example: Restoring wetland buffers that absorb flood shocks improves U7/D5.; Reviewer check: A reviewer asks whether proximity to thresholds was assessed, not just current condition.
U7/D6 - Life Continuity (label: life continuity)	Evidence: extinction-risk (e.g. red-list) data; biodiversity-trend indices; irreversibility assessments; Gate classes: Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Strongly tail-relevant: extinction is the paradigm of irreversible, non-averagable harm under TRC; containment-relevant.	Example: A habitat-corridor program that prevents a species' extinction improves U7/D6.; Reviewer check: A reviewer checks whether irreversibility was treated as a tail constraint, not a discountable cost.
U7/D7 - Ecological Integrity (label: ecological integrity)	Evidence: biodiversity indices; habitat-area data; water/soil quality; ecosystem-resilience measures; ecological monitoring; Gate classes: Rights Floor; Tail Risk; Containment / Structural Viability; Ordinary welfare. Uncertainty: Standard assessment. Interpretation: Often containment-relevant and sometimes catastrophe-relevant; may relate to ecological-integrity rights where recognised (e.g. rights of nature).	Example: A factory expansion that damages a watershed may fail containment even when it improves U4/D1.; Reviewer check: A reviewer checks whether integrity is evidenced by ecosystem-level data or inferred from a single metric.

APPENDIX AE: Carried-Forward v10.8 Layer-Discipline, Falsification, and Adaptive Reserve Addendum (Origin retained; active requirements governed by Appendix RELEASE)

Truth-surface discipline. Formal truth means that outputs follow from declared inputs and equations. Artifact truth means hashes identify the released artifacts. Implementation truth means validators reproduce expected gates, scores, and flags. Evidence truth means traces support or fail to support material claims. Empirical truth requires pilots, backtests, or real-world monitoring. Normative defensibility requires declared rights floors and ethical commitments that can be challenged, justified, and held. Governance legitimacy requires lawful authority, participation, review, contestability, and audit boundaries. These claim types support one another but MUST NOT be collapsed into one undifferentiated claim of truth.

AE.1 Purpose and release boundary

This v10.8 addendum hardens the Canon's category-collapse discipline without importing or depending on any external physics, ontology, or proprietary framework. It translates a general lesson from recent comparative review into MathGov-native audit rules: a value that can be written, computed, projected, or preferred is not automatically admissible, evidenced, authorized, or claimable.

Normative effect. The canonical cascade remains Reality Grounding → Rights Floor → TRC → CSV → RLS. Reality Grounding is added as the public claim-authority level, not as a compensatory scoring gate. Rights floors, TRC/CVaR, Containment, PLSS residual-only discipline, SGP boundaries, and Tier 4 ProofPack restrictions remain unchanged.

AE.2 COMPUTABLE_BUT_INADMISSIBLE audit state

Definition. `COMPUTABLE_BUT_INADMISSIBLE` is an audit state for cases where arithmetic, a formula, a model, or a workflow returns a value, but the declared domain rules prohibit treating that value as a valid decision state, conformance claim, or authorization basis.

Use cases include: a welfare score for an option that failed NCRC; an RLS ranking after TRC or CSV failure; a PLSS result used as if it could narrow admissibility; an SGP protection surface used beyond its evidence tier; a zero entered where the true state is unknown; or a reference vector that computes while violating declared closure boundaries.

Decision effect. A `COMPUTABLE_BUT_INADMISSIBLE` result may be logged, studied, and used for debugging. It **MUST NOT** authorize selection, deployment, conformance claims, or stronger maturity claims. The PCC **MUST** record why computation existed and why admissibility did not follow.

Alias note. `SGP_COMPUTABLE_BUT_INADMISSIBLE` and the `ripple.md` `ComputableButInadmissible` field are domain-scoped aliases of this state for SGP-specific and wrapper-specific records.

AE.3 Layer-specific falsification matrix

Layer	Layer-specific falsifier	Does not by itself falsify	Required response
NCRC	Rights-floor breach, unresolved rights uncertainty, or missing rights evidence for a rights-relevant case.	RLS arithmetic or UCI/HOI diagnostics.	Reject, refuse, or escalate under rights uncertainty.
TRC	Tail-risk/CVaR exceeds governed threshold or catastrophe scenario evidence is insufficient.	NCRC analysis already performed.	Reject or escalate; do not compensate with welfare score.
Containment	Local gain degrades containing systems beyond declared tolerance or hides cross-scale degradation.	Rights or tail-risk verdicts.	Reject as selectable; record containment rationale.
RLS	Wrong welfare cells, weights, evidence grade, aggregation, or sensitivity posture.	Admissibility gate outcomes.	Recompute or mark non-decisive; do not reopen failed gates.
UCI/HOI	Structural metric evidence, normalization, or tie-break application is defective.	NCRC/TRC/CSV/RLS by itself.	Recompute diagnostic or remove tie-break claim.
PCC / audit	Missing provenance, closure, evidence, reviewer trace, or claim-boundary record.	The underlying ethical truth of the situation.	Invalidate conformance or maturity claim until record is repaired.
SGP interface	SGP output lacks evidence tier, scoring version, reviewer status, uncertainty, or maturity boundary for its claimed use.	NCRC, TRC, CSV, lawful authority, or governance-role requirements.	Treat as governed evidence limitation; do not use as override.

AE.4 Reference, preference, number, and claim discipline

For RippleLogic runs, a numerical value is a representation of a declared evidence process, not truth itself. Implementations SHOULD distinguish: evidence source, measurement reference, governance preference or weighting choice, numerical representation, and public claim. Collapsing these layers is a category-collapse defect.

Reference/preference rule. A measurement reference identifies how a claim is grounded. A governance preference identifies an explicit normative or procedural choice, such as a weight, threshold, epsilon, or materiality rule. A number is the final representation after those choices. The PCC SHOULD disclose all three when the value is decision-material.

AE.5 Adaptive Reserve / Saturation Diagnostic

Adaptive Reserve is a carried-forward v10.8 structural diagnostic for detecting whether an option consumes so much social, institutional, ecological, cognitive, financial, or operational slack that the system becomes brittle even when immediate welfare scores appear positive.

Status. Adaptive Reserve is not a sixth gate in the current v13.0 line. It is a UCI/HOI-compatible diagnostic and containment-support signal. It MAY trigger review, sensitivity analysis, or escalation. It MUST NOT rescue an option that failed NCRC, TRC, or CSV.

Minimum fields. When used, implementations SHOULD record reserve domain, baseline reserve, projected reserve after option, recovery path, redundancy state, reversibility, monitoring trigger, and saturation risk.

AE.6 v10.8 release delta

Release Delta (v10.7 → v10.8). v10.8 adds explicit computable-but-inadmissible discipline, a layer-specific falsification matrix, a reference/preference/number claim boundary, and an optional Adaptive Reserve diagnostic. It does not alter the cascade, equations, rights thresholds, TRC mechanics, Containment semantics, SGP binding interface, PLSS residual-only rule, or Tier 4 claim boundary.

APPENDIX AF: Decision Verdict Semantics, Projection Freeze, Closure Status, and Authority-Selection Separation

AF.1 Purpose and release boundary

This v10.8 hardening appendix strengthens RippleLogic's own decision-state discipline without importing any external ontology, physics claim, proprietary framework, or hidden implementation layer. It preserves the canonical cascade Reality Grounding → Rights Floor → TRC → CSV → RLS. It does not change any rights threshold, TRC/CVaR rule, Containment rule, PLSS rule, SGP interface, UCI/HOI definition, or Tier-4 claim boundary.

Normative effect. Appendix AF adds output semantics, audit fields, and refusal/authority-separation rules. It clarifies what a RippleLogic run may claim when gates pass but the remaining admissible options are non-unique, evidence is insufficient, coverage is incomplete, or RLS/tie-breaks are non-decisive.

AF.2 Decision verdict states

FrameworkVerdict is the technical verdict emitted by the declared cascade, evidence-status checks, thresholds, scoring rules, tie-breaks, and refusal conditions. AuthoritySelection is a separate governance or operator act made after the FrameworkVerdict has been disclosed. Implementations MUST NOT present an authority selection as a deterministic framework selection.

The allowed FrameworkVerdict values are: BLOCK, REFUSE_DETERMINISTIC_SELECTION, ALLOW_FRAMEWORK_SELECTION, and ESCALATE. AUTHORIZED_SELECTION is not a

FrameworkVerdict; it is an authority record linked to the verdict when lawful or legitimate authority chooses among disclosed admissible options.

FrameworkVerdict / run-record crosswalk (Normative). ALLOW_FRAMEWORK_SELECTION serializes as SELECTED_DECISIVE only when the technical selection is unique and all selection conditions hold. REFUSE_DETERMINISTIC_SELECTION with a separately lawful, disclosed AuthoritySelectionRecord serializes as SELECTED_BY_AUTHORITY_NON_DECISIVE; without such a selection it serializes as REFUSE or DELAY according to the declared disposition. A controlled but not final execution posture serializes as PROVISIONAL_WITH_CONTROLS. An empty selectable set serializes as NO_SELECTABLE_OPTION. A governed TRC/CSV emergency-provisional posture serializes as EMERGENCY_PROVISIONAL. Redesign, evidence/authority escalation, and governed deferral serialize as REDESIGN, ESCALATE, and DELAY. BLOCK is a framework verdict, not a unique run-record state: its required serialization depends on the disclosed cause—normally NO_SELECTABLE_OPTION, REFUSE, REDESIGN, or ESCALATE. Implementations MUST record the cause rather than infer one state from the word BLOCK alone.

BLOCK means that no option may proceed under the claimed posture because a non-negotiable gate fails, no option remains admissible, or a required coverage/evidence condition invalidates the run.

REFUSE_DETERMINISTIC_SELECTION means that at least one option remains admissible, but the framework refuses to claim a unique technical selection because evidence is insufficient, coverage is incomplete, the option set is non-unique, RLS is non-decisive, or the Section 10.4/10.4B unique-selection conditions remain unmet. A governed tie-break preference does not remove this refusal (Section 11.5). ALLOW_FRAMEWORK_SELECTION requires complete qualification and the applicable unique-selection conditions under the stated tier, evidence and projection boundary. With multiple selectable options, Sections 10.4 and 10.4B govern; a non-decisive tie-break preference alone is insufficient. Sole-survivor handling remains as declared in Section 10.4, without a fabricated pairwise Gap. ESCALATE means additional evidence, higher-tier review, stakeholder challenge, expert review, or governance authorization is required.

AF.3 Non-decisive admissible sets

When NCRC, TRC, and required Containment checks leave multiple admissible options and the Section 10.4/10.4B every-contender and robustness conditions do not establish a decisive selection, RippleLogic MUST NOT claim unique deterministic framework selection. The correct technical outcome is REFUSE_DETERMINISTIC_SELECTION with an AdmissibleOptionSet and a reason code. A governance process MAY still choose among admissible options, but that choice MUST be recorded separately as an AuthoritySelectionRecord.

Minimum fields for a non-decisive admissible set are: FrameworkVerdict, AdmissibleOptionSet, DeterministicSelectionClaim=false, NonDecisiveReason, AuthoritySelectionRecord if one exists, AuthorityBasis, ReviewerRequired, and DisclosureText.

AF.4 ProjectionPreRegistration visibility

For Tier 3, high-stakes Tier 2, and future ProofPack-prep claims, implementations SHOULD expose whether a ProjectionPreRegistration object exists, what it hash-pins, and whether any material post-observation modification occurred. Scenario sets, CVaR confidence, TRC thresholds, uncertainty parameters, weights, scope declarations, SGP bindings, and refusal/escalation conditions SHOULD be frozen before the run for any claim-bearing evaluation. Post-observation modification creates a new run boundary unless explicitly recorded as exploratory analysis.

AF.5 Closure and coverage status

Required evidence surfaces MUST NOT silently default missing information to zero. Each required cell, right, scenario input, stakeholder mapping, SGP object, and projection input SHOULD carry a CoverageStatus: VALUE, ESTIMATE, ASSUMPTION, BLANK_BY_DESIGN, DECLARED_UNKNOWN, NOT_APPLICABLE, or MISSING_UNDECLARED. DECLARED_UNKNOWN may trigger REFUSE or ESCALATE when the unknown affects admissibility, selection, or tier claim. MISSING_UNDECLARED is non-conformant for claim-bearing runs.

AF.6 Falsification classes

Validation and failure records SHOULD identify the failure class rather than collapsing all failed tests into a generic defect. The recommended FalsificationClass values are DATA_FAILURE, PROJECTION_FAILURE, PARAMETER_FAILURE, OPERATOR_FAILURE, PRIMITIVE_CHALLENGE, GOVERNANCE_FAILURE, and COMMUNICATION_FAILURE. A failed projection or bad parameter does not automatically falsify the Canon; it identifies the layer that must be repaired or retested.

AF.7 Authority-selection record

AuthoritySelectionRecord documents a human, institutional, democratic, expert, emergency, or operator selection made after the framework has disclosed its technical verdict. It MUST identify selected option, authority type, authority basis, admissible option set, non-decisiveness disclosure, reviewer requirement, and whether the selection is lawful or policy-authorized within the deployment context. Authority selection does not convert REFUSE_DETERMINISTIC_SELECTION into ALLOW_FRAMEWORK_SELECTION.

AF.8 Reality-contact and primitive-clarity discipline

Reality-contact discipline clarifies P1 without expanding RippleLogic into a total metaphysics. A failed claim about reality does not mean reality failed; it means the evidence, projection, model, interpretation, or authority claim failed to make adequate contact with the relevant reality surface.

Reality-contact failures SHALL be classified for Tier 2 and Tier 3 claim-bearing runs as one or more of: DATA_FAILURE, PROJECTION_FAILURE, PARAMETER_FAILURE, OPERATOR_FAILURE, GOVERNANCE_FAILURE, COMMUNICATION_FAILURE, or PRIMITIVE_CHALLENGE. A revised claim after a reality-grounding failure is a new claim; it SHALL NOT be counted as the original claim having succeeded.

The framework itself is part of the social and institutional reality it affects. Any deployment of MathGov, RippleLogic, SGP, ripple.md, or RL_Agent that materially shapes decisions, incentives, authority, or public belief is subject to the same reality-grounding audit discipline it applies to other systems.

AF.9 v10.8 release delta

Release Delta (v10.6 → v10.8). The v10.8 release line carries forward decision-verdict semantics, authority-selection separation, projection-freeze visibility, coverage-status closure discipline, falsification-class taxonomy, and the root reality doctrine / reality-grounding consistency patch. It is a hardening patch only. It does not add a new gate, alter the cascade, import an external ontology, or make Tier 4, ProofPack, empirical-validation, legal-certification, or deployment-readiness claims available.

Appendices AG and AH carry the v10.8 Reality Grounding / Lean Cascade and Category Grounding patches and are part of this canonical artifact.

APPENDIX AG: CARRIED-FORWARD v10.8 REALITY GROUNDING AND LEAN CASCADE PATCH (CURRENTLY PRESERVED THROUGH v13.0)

Status. Appendix AG records the v10.8 substance update carried forward through the v13.0 release line. Appendix AC remains controlling for release metadata and claim boundaries where conflicts appear.

Core method. The public-facing RippleLogic method inside the MathGov framework is now: Reality Grounding -> Rights Floor -> TRC -> CSV -> RLS, or compactly RG -> RF -> TRC -> CSV -> RLS.

Formal method. The formal implementation is: RG/RSG -> RF/NCRC -> TRC -> CSV -> RLS. RF names the public layer; NCRC names the formal rights predicate.

Reality Grounding. Level 1 declares the reality surface, evidence trace, material unknowns, transition boundary, consequence pathways, and claim boundary. No claim may exceed its reality surface.

UCI/HOI placement. UCI/HOI are structural diagnostic instruments, not public headline cascade stages. Gate-relevant UCI/HOI evidence belongs inside CSV. Residual UCI/HOI applies only when RLS is tied, close, uncertainty-overlapped, or non-decisive, or when monitoring/hollowing documentation is required.

PLSS. PLSS remains the zoom lens, not the safety waiver. It adjusts residual welfare-ranking attention only after admissibility has been preserved and cannot omit containment-relevant structural diagnostics.

Boundary. This release does not claim Tier 4, ProofPack readiness, empirical validation, legal certification, deployment certification, or deterministic selection where the run is non-decisive.

APPENDIX AH: CATEGORY GROUNDING AND PRIMITIVE AUDIT PROTOCOL

AH.1 Status and placement

Status. Appendix AH is a normative clarification for claim-bearing Tier 2 and Tier 3 runs and an informative design guide for lower-stakes runs. It is part of the Reality Grounding layer. It does not create a sixth public method level and does not weaken NCRC, TRC, CSV, RLS, PLSS boundaries, SGP boundaries, lawful authority boundaries, or Tier 4/ProofPack boundaries.

Placement rule. Category Grounding, Primitive Audit, Reference Category Review, Definition Falsification Test, and Category-Reality Correspondence Check are one grouped subdiscipline under Reality Grounding when material. Implementers MAY use these names as subfields or review views, but MUST NOT split them into competing cascade stages or treat them as optional decorative terminology when category choice is material.

AH.2 Definitions

Category. A category is an operational class used by a RippleLogic run to identify, group, measure, protect, score, govern, authorize, or execute against phenomena. Examples include person, stakeholder, union scope, harm, right, sentience, agency, intelligence, safety, stability, risk, authority, protection, flourishing, containment, and alignment.

Category Grounding. Category Grounding is the Reality Grounding subdiscipline that checks whether a material category has a declared definition, inclusion boundary, exclusion boundary, reference structure, evidence surface, maturity status, and falsification or revision trigger sufficient for the claim being made.

Primitive Audit. Primitive Audit is the review of a framework starting point or imported foundational term, including its dependency basis, operational role, limits, plausible failure modes, and revision conditions.

Primitive Audit is required when a run, release, or public claim materially depends on a primitive in a contested or high-consequence way.

Semantic Debt. Semantic debt is the accumulated governance, measurement, operational, and ethical burden created when a decision system continues to use categories whose reality correspondence, boundary conditions, or falsification status are weak, stale, inherited, contested, or unreviewed. Semantic debt is not automatically a gate failure, but it raises review priority and may require claim narrowing, monitoring, escalation, or refusal when material.

Term Discrimination. Term Discrimination is the Category Grounding subdiscipline that states what a material term means, what it does not mean, what discriminator separates it from nearby terms, where it sits in the dependency order, and when the term must be narrowed, escalated, monitored, treated as sensitivity-only, or refused.

Category Discriminator. A category discriminator is the decisive property that separates a material category or term from neighboring categories it is likely to be confused with.

Dependency Position. Dependency position is the location of the term in the decision stack, such as reality-surface, evidence-trace, category-boundary, rights-admissibility, tail-risk-admissibility, containment-selectability, residual-ranking, framework-verdict, authority-selection, execution-authority, monitoring, certification, or public-claim.

AH.3 Trigger conditions

A CategoryGroundingRecord is REQUIRED for claim-bearing Tier 3 runs, high-stakes Tier 2 runs, and public conformance claims when any of the following are material: (1) the category affects NCRC protection, stakeholder inclusion, or rights-floor localization; (2) the category affects TRC scenario construction, tail probabilities, severity bands, or irreversibility classification; (3) the category affects Containment, UCI, HOI, Adaptive Reserve, lock-in, hollowing, domination, institutional erosion, ecological degradation, or governance-capacity loss; (4) the category affects SGP outputs, moral-patient status, protection, standing, or governance-authority interpretation; (5) the category is a metric target under strong optimization pressure; (6) the category is novel, inherited, contested, ambiguous, or cross-domain; (7) the category supports an authority-selection, deployment, certification, validation, or public maturity claim.

AH.4 Minimum CategoryGroundingRecord

Minimum fields: material_category; category_role; primitive_basis; definition; category_discriminator; dependency_position; inclusion_boundary; exclusion_boundary; common_confusions; noncollapse_pairs; reference_structure; evidence_surface; source_authority_status; category_maturity_status; falsification_or_revision_trigger; term_refusal_condition; downstream_dependencies; semantic_debt_flag; reviewer_status; and required_claim_action.

Allowed category_maturity_status values: GROUNDED, ESTIMATED, ASSUMPTION_BOUND, CONTESTED, LEGACY_UNREVIEWED, INSUFFICIENT_GROUNDING, and OUT_OF_SCOPE_WITH_RATIONALE.

PrimitiveAuditRecord SHOULD include: primitive_id; primitive_statement; primitive_role; imported_parent_disciplines; operational_dependency; affected_claims; evidence_or_argument_basis; failure_mode; falsification_or_revision_trigger; reviewer_status; required_claim_action; and audit_reference. PrimitiveAuditRecord is required when a primitive or imported foundational term materially supports a rights, tail-risk, containment, SGP, authority, conformance, public-claim, or deterministic-selection claim.

Allowed semantic_debt_flag values: NONE, LOW, MEDIUM, HIGH, and CRITICAL. HIGH or CRITICAL semantic debt requires explicit reviewer attention before stronger public, conformance, deployment, authority, or deterministic-selection claims may be made.

Semantic debt rubric: NONE = category recently reviewed, evidence-linked, and uncontested within the claim boundary. LOW = minor ambiguity with no plausible gate or claim effect. MEDIUM = ambiguity could affect

interpretation or monitoring but not current admissibility. HIGH = ambiguity could affect rights, tail risk, containment, authority, conformance, or public claim. CRITICAL = category is stale, contested, or ungrounded in a way that could invalidate the decision state or public claim.

Allowed `required_claim_action` values: ALLOW, NARROW, ESCALATE, REFUSE, MONITOR_ONLY, and SENSITIVITY_ONLY.

Allowed `source_authority_status` values SHOULD include: PRIMARY_EVIDENCE, GOVERNING_SPECIFICATION, DOMAIN_STANDARD, LEGAL_AUTHORITY, INSTITUTIONAL_POLICY, EXPERT_JUDGMENT, MODEL_OUTPUT, STAKEHOLDER_TESTIMONY, LEGACY_SOURCE, and UNVERIFIED_SOURCE. A source authority label records provenance; it does not by itself create admissibility, validity, legality, or execution authority.

Allowed `dependency_position` values SHOULD include: REALITY_SURFACE, EVIDENCE_TRACE, CATEGORY_BOUNDARY, RIGHTS_ADMISSIBILITY, TAIL_RISK_ADMISSIBILITY, CONTAINMENT_SELECTABILITY, RESIDUAL_RANKING, FRAMEWORK_VERDICT, AUTHORITY_SELECTION, EXECUTION_AUTHORITY, MONITORING, CERTIFICATION, and PUBLIC_CLAIM.

AH.4A Aggregate Semantic-Debt Ledger (Normative for material claim-bearing runs)

Per-category semantic debt may accumulate into a run-level liability even when no single category is independently fatal. Claim-bearing Tier 2 and Tier 3 runs SHOULD compute a `semantic_debt_score` over material categories using declared weights: NONE=0, LOW=1, MEDIUM=3, HIGH=7, CRITICAL=15.

If any material category is CRITICAL, stronger public, conformance, authority-selection, deployment-readiness, or deterministic-selection claims MUST be refused, narrowed, or escalated unless a reviewer records a governed disposition. If `semantic_debt_score` $\geq$ `threshold_T`, with `threshold_T` declared ex ante in the verification regime and default `threshold_T` = 10 for worked examples, implementations SHOULD flag `SEMANTIC_DEBT_AGGREGATE_HIGH` with severity ESCALATE and require reviewer disposition before stronger public or conformance claims.

Semantic-debt aggregation is not an ontology claim. It is an auditability rule for accumulated category uncertainty inside Reality Grounding.

AH.5 Category failure classes

Category-relevant failures SHOULD be tagged as one or more of: CATEGORY_BOUNDARY_FAILURE, CATEGORY_REFERENCE_FAILURE, CATEGORY_SCOPE_FAILURE, CATEGORY_METRIC_MISMATCH, SEMANTIC_DRIFT_DETECTED, TERM_DISCRIMINATION_FAILURE, DEPENDENCY_POSITION_ERROR, PRIMITIVE_CHALLENGE, DATA_FAILURE, PROJECTION_FAILURE, PARAMETER_FAILURE, OPERATOR_FAILURE, GOVERNANCE_FAILURE, or EXTERNAL_CHANGED_CONDITION.

CATEGORY_BOUNDARY_FAILURE means the inclusion or exclusion rule for a material category is wrong, missing, or too vague for the claim. CATEGORY_REFERENCE_FAILURE means the category lacks a valid reference structure. CATEGORY_SCOPE_FAILURE means the category was applied at the wrong union scope, stakeholder class, substrate, time horizon, or decision boundary. CATEGORY_METRIC_MISMATCH means the metric does not track the category it claims to measure. SEMANTIC_DRIFT_DETECTED means ordinary, institutional, legal, benchmark, or model usage has drifted from the category structure required for the run. PRIMITIVE_CHALLENGE means the framework starting point itself may require revision, narrowing, or retirement.

Category-failure action rule. If a material category is CATEGORY_BOUNDARY_FAILURE, CATEGORY_REFERENCE_FAILURE, CATEGORY_METRIC_MISMATCH, SEMANTIC_DRIFT_DETECTED at HIGH/CRITICAL severity, LEGACY_UNREVIEWED for a claim-bearing run, or otherwise insufficiently grounded for the claimed decision state, the run SHALL NOT produce an unqualified deterministic framework

selection. It MUST narrow, escalate, mark sensitivity-only or monitor-only, collect stronger evidence, or refuse the stronger claim. A computable score over a category-invalid surface is computable but inadmissible.

AH.6 Relationship to optimization, governance, and public claims

Optimization is admissible only after the relevant category, evidence surface, and claim boundary are adequate for the intended claim. Governance, compliance, benchmark success, institutional permission, model fluency, dashboard availability, or stakeholder preference does not by itself validate a category. If the category is materially wrong, a higher RLS score, a completed PCC, a PLSS profile, or an authority-selection record cannot repair it.

AH.7 Release boundary

This appendix does not claim empirical validation, legal certification, deployment certification, deterministic selection, or Tier 4/ProofPack readiness. It strengthens the category-discipline portion of Reality Grounding so that MathGov does not become a governance layer over unexamined definitions. This carries forward the Category Grounding discipline introduced in RippleLogic Framework v12.4 into the v12.6 line, including Sections 1 through 20 and Appendices A through AH; Appendix AI then adds the Term Discrimination and Semantic Stability Protocol as a normative subdiscipline of Category Grounding where material.

McGaughan, J. (2026). Sentience Gradient Protocol v8.6: A theory-plural, evidence-governed framework for moral patienthood, protection, participation, stewardship, and reality-management calibration across substrates. MathGov Core Release 2026.09 v12.7.

APPENDIX AI: TERM DISCRIMINATION AND SEMANTIC STABILITY PROTOCOL

AI.1 Status and placement

Status. Appendix AI is a normative hardening patch for claim-bearing Tier 2 and Tier 3 runs and an informative design guide for lower-stakes runs. It strengthens Category Grounding inside Reality Grounding. It does not create a sixth public method level, does not import an external ontology, does not weaken NCRC, TRC, CSV, RLS, PLSS, SGP, lawful authority, or Tier 4/ProofPack boundaries, and does not make certification, compliance, or deployment claims available by itself.

Placement rule. Term Discrimination is a required subdiscipline of Category Grounding when a material term affects rights, tail risk, containment, SGP interpretation, authority, conformance, public claims, optimized metrics, or execution. Category Grounding asks whether a category is grounded. Term Discrimination asks whether the operative word is being kept separate from neighboring words that would change the claim, gate result, authority state, or execution status if collapsed.

AI.2 Core definitions

Term. A term is an operative word, phrase, label, token, schema field, policy label, benchmark label, dashboard label, legal label, or model-generated description used to identify, measure, protect, score, govern, authorize, certify, monitor, or execute against a phenomenon.

Term Discrimination. Term Discrimination is the review discipline that states what a material term means, what it does not mean, what discriminator separates it from neighboring terms, where it belongs in the dependency order of the run, and when the term must be narrowed, escalated, monitored, or refused.

Discriminator. A discriminator is the decisive property that makes one term this term rather than a neighboring term. A discriminator MUST be stated for material terms that could otherwise collapse into another term under institutional pressure, optimization pressure, model fluency, compliance theater, or public communication pressure.

Dependency position. Dependency position identifies where a term sits in the decision stack. A downstream term **MUST NOT** claim the authority of an upstream term. Permission does not establish admissibility. Monitoring does not establish control. Certification does not establish validity. A score does not establish truth. Execution capability does not establish execution authority.

Term-refusal condition. A term-refusal condition is the condition under which a run **MUST NOT** use a term at its stronger claim level. The required response is claim narrowing, evidence collection, escalation, sensitivity-only use, monitor-only use, or refusal.

Non-collapse pair. A non-collapse pair is a pair of terms that **MUST NOT** be treated as equivalent unless a future governed annex explicitly defines an equivalence for the relevant use. Non-collapse pairs include reality/evidence, evidence/interpretation, compliance/correctness, certification/validity, permission/admissibility, monitoring/control, safety/stability, automated execution/operational autonomy, capability/authority, protection/authority, selected/executable, and score/truth. The distinction does not prohibit bounded functional use of capability terms; it prohibits using a weaker observation as proof of a stronger claim.

Autonomy descriptor. Autonomy is not treated as a single ordered capability. A material autonomy claim is a vector describing the operational definition, observable criteria, system composition, evidence surface, validity conditions, alternative explanations, candidate generation, constraints or references, objective source and change authority, admissibility warrants, execution interface and scope, authorization, revocation, safe state, falsification trigger, evidence maturity, permitted claim, and explicit nonclaims. A bounded system-level capability may be real even when implemented through ordinary components; no component, label, or benchmark supplies the evidence or authority of the others.

AI.3 Extended CategoryGroundingRecord fields

When Term Discrimination is material, the CategoryGroundingRecord **MUST** add or map the following fields:

Field	Meaning
category_discriminator	The property that separates the material category or term from nearby terms it is likely to be confused with.
dependency_position	The layer or authority position the term occupies: reality-surface, evidence-trace, category-boundary, rights-admissibility, tail-risk-admissibility, containment-selectability, residual-ranking, framework-verdict, authority-selection, execution-authority, monitoring, certification, or public-claim.
term_refusal_condition	The condition under which the term cannot support the stronger claim and must be narrowed, escalated, marked sensitivity-only/monitor-only, or refused.
noncollapse_pairs	The nearest terms this term MUST NOT be collapsed with in the run.
source_authority_status	Whether the source of the term is evidential, legal, institutional, model-generated, benchmark-derived, stakeholder-proposed, or framework-defined. This records source status without treating source status as automatic category validity.
capability_authority_decomposition	For material automation/autonomy/agentic/reasoning claims: operational_definition; observable_criteria; mechanism_or_system_composition; evidence_surface; validity_conditions; alternative_explanations; falsification_or_revision_trigger; evidence_maturity_stage; permitted_claim; explicit_nonclaims; candidate_generation_method; reference_source; constraint_sources; objective_source; objective_change_authority; admissibility_warrant_sources; execution_interface; execution_scope; execution_authority_source; revocation_authority; safe_state_or_shutdown_path; and requalification_trigger.

AI.4 Dependency-position rule

A material term may support only the claim strength allowed by its dependency position and evidence surface.

Term position	May support	Must not claim by itself
Reality surface	Contact with the modeled or evidenced part of the territory	Complete access to the territory
Evidence trace	Support for a claim boundary	Self-interpreting truth or complete causality
Category boundary	Proper classification for declared use	Validity of downstream scoring or governance by itself
Rights-admissibility	NCRC pass/fail state	Tail-risk pass, containment pass, selection, or execution authority
Tail-risk-admissibility	TRC pass/fail state	Rights validity, containment integrity, or welfare optimality
Containment-selectability	Larger-system structural acceptability	Highest residual welfare or lawful execution
Residual-ranking	RLS ranking among selectable options	Rescue of any failed upstream layer
Framework verdict	Disclosed result of the framework run	Human, legal, institutional, or execution authority
Authority selection	A human/institutional/operator choice after verdict disclosure	Conversion of non-decisiveness into deterministic framework selection
Execution authority	Governed authority to act through hard controls	Proof that the action is safe, valid, or wise
Monitoring	Observability and feedback after or during action	Prevention, control, safety, or stability by itself
Certification	Evidence of conformance to a declared standard	Reality validity, legal compliance, empirical truth, deployment readiness, or Tier 4 status

AI.5 Do-Not-Collapse Matrix

Do not collapse	Discriminator	Required MathGov response
Reality vs evidence	Reality is the evaluator-independent constraint-and-consequence field; evidence is traceable contact with a surface of it.	Keep claim boundary inside the evidence surface.
Evidence vs interpretation	Evidence is trace; interpretation connects trace, uncertainty, values, and action.	Record reviewer status and uncertainty.
Category vs metric	Category defines what is being measured; metric quantifies a proxy or indicator.	Refuse or narrow metric claims when the category is weak.
Score vs truth	RLS ranks selectable options under declared inputs; it is not truth itself.	Never let RLS rescue failed grounding, rights, risk, or containment.
Compliance vs correctness	Compliance means conformity to a rule or process; correctness means reality-bound adequacy for the claim.	Treat compliance as evidence, not validity.
Certification vs validity	Certification records an assurance result; validity depends on grounded reference, evidence, and replayable adequacy.	Prevent certification theater and overclaiming.
Permission vs admissibility	Permission is authorization; admissibility is passing the relevant MathGov gates.	Do not treat approval as gate success.
Monitoring vs control	Monitoring observes; control prevents, constrains, or corrects.	Require enforcement path, rollback, or containment evidence for control claims.

Do not collapse	Discriminator	Required MathGov response
Safety vs stability	Safety means bounded non-harm under declared conditions; stability means preserved structural regime integrity.	Use TRC, Containment, and Adaptive Reserve where relevant.
Stability vs justice	Stability can preserve harmful arrangements; justice requires rights floors and non-domination.	Do not allow stable oppression to pass rights review.
Automated execution vs operational autonomy	Automation alone does not establish autonomy, but coordinated memory, planning, feedback, tool use, and bounded execution may instantiate a system-level operational capability when defined and demonstrated.	Require an operational definition, observable criteria, system composition, alternatives, falsifier, evidence maturity, and the full capability-authority decomposition. Do not infer consciousness, moral agency, generality, warrant, or authority from the operational capability.
Probabilistic inference vs execution warrant	Probabilistic inference can generate, estimate, classify, or propose; an execution warrant supports a bounded transition through domain evidence and declared controls.	Use inference as evidence only within its validated role; require PC-AEP, CSV, authority, and controls where consequence-bearing execution is material.
Protection vs authority	Protection concerns moral/welfare standing; authority concerns legitimate role, responsibility, competence, and authorization.	Preserve SGP protection/authority separation.
Selected vs executable	A selected option is chosen by framework or authority; executable means action may proceed through lawful/runtime controls.	Require execution-stability and authority checks.
Capability vs authority	A system may be technically able to act without being permitted or admissible to act.	Block action without hard-control authorization.

AI.6 Claim-action rule

If a material term lacks a discriminator, lacks a dependency position, collapses into a neighboring term, or lacks a refusal condition sufficient for the claim being made, the run MUST classify the issue as TERM_DISCRIMINATION_FAILURE or a more specific category failure class. The result MUST be NARROW, ESCALATE, REFUSE, MONITOR_ONLY, or SENSITIVITY_ONLY until the term is grounded. A completed PCC, authority-selection record, institutional permission, model output, dashboard result, or certification label MUST NOT repair a term-discrimination failure by itself.

AI.7 v12.0 release delta

Release Delta (v11.0 -> v12.0). v12.0 carries forward the Term Discrimination and Semantic Stability Protocol and adds Structural Viability and Falsification Discipline hardening inside existing layers. The v11.0 line added Term Discrimination as a hardening patch inside Category Grounding and Reality Grounding. It adds category_discriminator, dependency_position, term_refusal_condition, noncollapse_pairs, and source_authority_status to material category records; adds the Do-Not-Collapse Matrix; strengthens certification/permission/monitoring non-authority language; and preserves the five-level method. It does not add a new public gate, import an external ontology, alter NCRC, TRC, CSV, RLS, SGP protection surface semantics, PLSS, lawful authority boundaries, Tier 4, ProofPack, empirical-validation, legal-certification, or deployment-certification boundaries.

Appendix AJ: Public Release Hardening Addendum (v12.0)

Status. This appendix is a hardening bridge for the v12.0 content line. It does not create Tier 4, ProofPack, legal certification, deployment certification, empirical validation, or automated moral truth. It clarifies implementation discipline and points to companion guidance. If a companion guide conflicts with the Canon, the Canon controls until a future versioned amendment resolves the conflict.

AJ.1 Gate Boundary Discriminator: TRC, CSV, and RLS

TRC, CSV, and RLS SHALL NOT be collapsed.

Layer	Primary question	Normal consequence
TRC	Could this option create catastrophic, irreversible, ruinous, or lock-in tail exposure?	Gate failure if tail exposure exceeds corridor.
CSV	Can the option deliver its claimed good without breaking, hollowing, overloading, exploiting, or depending on containing systems beyond tolerance?	Fail, controls, redesign, escalation, or pass depending on structural status.
RLS	Among options that survived RG, RF/NCRC, TRC, and CSV, which has the strongest residual welfare profile?	Ranking only; no gate rescue.

Double-materiality rule. A harm may be both tail-relevant and structurally degrading. In that case it **MUST** be evaluated in TRC and CSV as appropriate, recorded in the PCC, and not double-counted as an ordinary residual RLS cost unless a distinct residual effect remains after gate treatment.

Companion guide: docs/guides/GATE_BOUNDARY_DISCRIMINATOR_TRC_CSV_RLS.md.

AJ.2 Rights Threshold Governance

Rights thresholds are governance commitments under evidence, not preferences to be moved after selection. Tier 3 and public claim-bearing Tier 2 runs **SHOULD** record threshold source, rationale, under-protection risk, over-blocking risk, review status, and revision trigger. A decision owner who benefits from lowering a threshold **MUST NOT** unilaterally lower the Rights Floor for the run under evaluation.

Companion note: docs/validation/rights/RIGHTS_THRESHOLD_GOVERNANCE_NOTE_v1_0.md.

AJ.3 Subgroup Discovery and Worst-Affected Groups

Worst-off subgroup semantics require active discovery of groups likely to bear hidden, indirect, intersectional, or least-contestable harms. Unknown subgroup exposure is not zero. Where direct sensitive-data collection is unsafe or unlawful, pathway analysis, proxy-risk statements, consultation, conservative bounds, and challenger review **SHOULD** be used.

Companion guide: docs/guides/SUBGROUP_DISCOVERY_AND_WORST_AFFECTED_GROUPS.md.

AJ.4 Emergency Mode Least-Rights-Infringing Rule

Emergency Mode does not make a rights-failing option aligned. It identifies a time-bounded, least-rights-infringing provisional action when no normal option passes RF/NCRC. RLS may be used only as a final constrained tie-break after rights-violation vector comparison and TRC comparison. RLS **MUST NOT** reverse the emergency rights-vector ordering or convert a rights-failing option into an ordinary selectable option.

Companion guide: docs/guides/EMERGENCY_MODE_LEAST_RIGHTS_INFRINGING_PROTOCOL.md.

AJ.5 TRC Scenario Discovery

TRC is only as complete as the declared scenario surface. Tier 3 runs **SHOULD** use scenario discovery sources including domain experts, incident and near-miss data, red teams, affected stakeholders, literature, standards, and AI-assisted candidate generation. AI-generated scenarios are discovery inputs, not validity sources. Material rejected scenarios **SHOULD** be logged with rationale.

Companion protocol: docs/validation/trc/TRC_SCENARIO_DISCOVERY_PROTOCOL_v1_0.md.

AJ.6 CSV Measurement Maturity

CSV claims SHOULD declare maturity when used in public or high-stakes contexts: qualitative, indicator-supported, calibrated, or validated. UCI/HOI-style diagnostics are structural signals and MUST NOT be presented as empirically validated in a domain unless domain validation exists.

Companion note: docs/validation/csv/CSV_MEASUREMENT_MATURITY_NOTE_v1_0.md.

AJ.7 PCC Profiles

PCC burden SHOULD be proportional to stakes without allowing consequential decisions to evade audit. The recognized profiles are PCC-Lite, PCC-Core, PCC-Audit, PCC-Agent, and PCC-Public. A lower profile is not allowed when rights exposure, catastrophe relevance, CSV burden, automation, public claims, affected-party burden, or accumulated systemic effect require a higher profile.

Companion guide: docs/guides/PCC_PROFILE_GUIDE.md.

AJ.8 HDW Participation and Capture Control

HDW democratic tuning SHOULD record affected groups, representation method, exclusion risks, capture risks, minority reports, and correction actions in Tier 3 runs. Optional mechanisms such as quadratic voting or deliberative polling MAY be used where appropriate, but no method may lower constitutional floors, erase non-vocal stakeholders, or weaken RF/NCRC, TRC, or CSV.

Companion guide: docs/guides/HDW_PARTICIPATION_AND_CAPTURE_GUIDE.md.

AJ.9 SGP Validation Bridge

SGP outputs remain protection-surface evidence objects, not legal status, governance authority, or substitutes for RG, RF/NCRC, TRC, CSV, RLS, lawful authority, or governance-role requirements. SGP validation should test inter-rater reliability, fluency traps, biological welfare false negatives, artificial-entity status inflation, and protection-authority separation.

Companion protocol: docs/validation/sgp/SGP_VALIDATION_PROTOCOL_v2_1.md.

AJ.10 Stewardship Evasion Audit

A run cannot avoid stewardship obligations by labeling consequential influence as educational, internal, experimental, advisory, or informational. If the output materially influences real decisions, narrows options, affects rights or livelihoods, or connects to execution capability, stewardship review SHOULD be escalated.

Companion guide: docs/guides/STEWARDSHIP_EVASION_AUDIT.md.

AJ.11 Meta-Union Firewall

AIU, Cosmic, Universal, and related horizon/meta-union language may orient humility and philosophical interpretation. They are not Tier 1-3 scoring objects, gate overrides, hidden weights, legal authority, deployment permission, or ProofPack evidence unless a future governed extension explicitly defines a computable and validated protocol.

Companion guide: docs/guides/META_UNION_FIREWALL.md.

Appendix AK: Computability vs Realizability Bridge

Status. Appendix AK is a release-hardening bridge introduced in v12.0 and explicitly precedence-indexed in the v12.0 final core. It does not add a new public cascade level and does not change the governing order RG -> RF -> TRC -> CSV -> RLS.

AK.1 Core distinction. A state, output, plan, policy, or model response may be computable without being physically grounded, rights-safe, ruin-bounded, structurally viable, authorized, or ethically selectable. Computability is representational capacity. Selectability is governed survivorship through the cascade.

AK.2 Operational rule. RippleLogic MUST NOT treat model fluency, simulation output, dashboard value, code generation, tool result, or computed score as decision authority merely because it exists. Computable is not selectable. Generated is not grounded. Executable is not ethical. Realizable is necessary but not sufficient.

AK.3 Layer assignment. RG handles reality-surface contact and claim boundaries. RF/NCRC handles non-compensatory rights protection. TRC handles catastrophic and ruin-path exposure. CSV handles containment, structural viability, execution viability, monitoring, authority dependencies, and host-system integrity. RLS ranks only residual welfare for options that remain selectable.

AK.4 Feedback rule. If CSV analysis discovers a new catastrophic, irreversible, lock-in, or ruin-path scenario that was not represented in TRC, the run MUST reopen TRC before RLS. CSV can discover TRC-relevant material, but it does not absorb or replace TRC.

AK.5 Public teaching line. Ground reality. Protect rights. Bound ruin. Preserve the structure. Score the ripples.

Appendix AL: Source-Coupling Integrity and Generative-Source Traceability (Normative Hardening Patch)

Status. Appendix AL is a v12.0 hardening patch. It is a MathGov-native Reality Grounding and CSV diagnostic. It does not add a public gate, does not alter the cascade, does not import an external ontology, and does not make empirical validation, legal certification, deployment certification, Tier 4, ProofPack, or automated moral-truth claims available.

AL.1 Purpose

Appendix AL protects RippleLogic runs from a specific downstream-substitution failure: treating output, fluency, institutional permission, compliance status, inherited procedure, benchmark performance, dashboard value, or interface success as proof that the claimed capability is grounded in the enabling conditions that make it possible and the boundary conditions that limit it.

AL.2 Required Source-Coupling Record

When the Section 2.1D trigger holds, the PCC SHOULD include a SourceCouplingRecord and MUST include one for Tier 3 public or institutional claim-bearing runs. The record contains: `claimed_capability`; `enabling_conditions`; `boundary_conditions`; `source_evidence`; `generator_output_distinction`; `inherited_assumptions`; `downstream_compensations`; `source_coupling_status`; `source_debt_flag`; `falsification_or_recheck_trigger`; and `required_claim_action`.

AL.3 Claim Boundary Effects

Source-Coupling Integrity constrains claim strength. A run MAY proceed with a narrow claim when source coupling is partial or inferred, but it MUST NOT use partial or inferred source coupling to support stronger conformance, deployment, deterministic-selection, or validation claims. If source coupling is unknown, contested, or failing, the run MUST narrow, collect evidence, escalate, redesign, or refuse the stronger claim.

AL.4 CSV Consumption Rule

If source weakness materially affects dependency closure, resource closure, operational capacity, reversibility, containment, monitoring, or host-system integrity, CSV MUST consume the SourceCouplingRecord. `SOURCE_DEBT_RISK` is a CSV-relevant signal. It may route the option to `CSV_PASS_WITH_CONTROLS`, `CSV_REDESIGN_REQUIRED`, `CSV_ESCALATE`, `CSV_FAIL`, or emergency-provisional handling according to the CSV Gate Standard and the Canon.

AL.5 Agentic and AI System Rule

For AI or agentic systems, model fluency, benchmark performance, tool-use success, chain-of-action success, or safety-filter compliance is not source-coupled admissibility. Such evidence may support a bounded capability claim, but only after Reality Grounding states what the evidence does and does not show, and only after CSV checks whether deployment can structurally stand.

AL.6 Related External Work Boundary

This Appendix is stated entirely in MathGov-native terminology as Source-Coupling Integrity. It checks whether a decision claim remains connected to its declared source, evidence surface, validity domain, and correction pathway. No external ontology, protected framework, or third-party text is incorporated into the governing Canon.

Appendix AM: Physical/Causal Admissibility Evidence Profile (Normative Hardening Patch)

Status. Appendix AM is a v12.0 Governed Admissibility hardening patch. It is a MathGov-native Reality Grounding and CSV profile for material physical or causal action. It does not add a public gate, does not alter the cascade, does not import an external ontology, and does not make empirical validation, legal certification, deployment certification, formal-methods proof, domain engineering certification, Tier 4, ProofPack, or automated moral-truth claims available.

AM.1 Purpose

Appendix AM protects RippleLogic runs from a specific execution-substitution failure: treating computed possibility, generated action, simulation output, policy approval, certification, compliance, monitoring, or governance permission as proof that a candidate physical or causal action is admissible before execution.

AM.2 Required Physical/Causal Admissibility Evidence Profile

When the Section 2.1E trigger holds, the PCC SHOULD include a PhysicalCausalAdmissibilityEvidenceProfile and MUST include one for Tier 3 public or institutional claim-bearing runs and high-stakes Tier 2 runs where physical or causal uncertainty could change gate status, authority, or public claim strength. The record contains: candidate_generation_source; physical_or_causal_model_used; validity_domain; boundary_conditions; uncertainty_range; failure_modes; reversibility_or_irreversibility_boundary; verification_simulation_empirical_test_or_expert_warrant; admissibility_warrant_source; monitoring_and_shutoff_path; residual_unknowns; and required_claim_action.

AM.3 Claim Boundary Effects

Physical/Causal Admissibility Evidence constrains claim strength. A run MAY proceed with a narrow claim when the evidence is partial, assumption-bound, or controlled, but it MUST NOT use partial physical or causal evidence to support stronger conformance, deployment, deterministic-selection, safety, reliability, or validation claims. If the model basis, validity domain, boundary conditions, uncertainty, failure modes, reversibility boundary, verification or warrant, monitoring/shutoff path, or residual unknowns are insufficient for the requested action, the run MUST narrow, control, redesign, escalate, or refuse.

AM.4 CSV Consumption Rule

CSV MUST consume the profile when physical or causal weakness materially affects dependency closure, resource closure, operational capacity, reversibility, monitoring adequacy, containment, execution viability, authority, or host-system integrity. The profile may route the option to CSV_PASS_WITH_CONTROLS, CSV_REDESIGN_REQUIRED, CSV_ESCALATE, CSV_FAIL, or emergency-provisional handling according to severity, reversibility, uncertainty, monitoring, authority, and available alternatives.

AM.5 TRC Reopen Rule

If the profile reveals a catastrophic, irreversible, lock-in, ruin-path, or severe-harm scenario not already represented in TRC, the run MUST reopen TRC before RLS. CSV can discover TRC-relevant material, but it does not absorb or replace TRC.

AM.6 AI, Agentic, and Physical-System Rule

For AI, robotics, cyber-physical, infrastructure, medical, industrial, environmental, or other consequence-bearing systems, model fluency, benchmark performance, tool-use success, chain-of-action success, safety-filter compliance, operator approval, certification, or monitoring is not physical or causal admissibility. Such evidence may support a bounded claim only after Reality Grounding states what the evidence does and does not show, the Physical/Causal Admissibility Evidence Profile declares the model and limits, and CSV checks whether the pathway can structurally stand.

AM.7 Boundary Against Overclaim

This profile is an evidence discipline, not a proof engine. It does not make MathGov a substitute for physics, control engineering, formal verification, medicine, ecology, cybersecurity, structural engineering, law, or domain expert review. Its function is to prevent a governance record from hiding the absence, weakness, uncertainty, or contested status of the physical or causal evidence on which execution depends.

Appendix AN: Methodological Falsifiability and Dependency Integrity Standard (Normative Hardening Patch)

Status. Appendix AN is a v12.0 Methodological Integrity hardening patch. It is a MathGov-native Reality Grounding discipline for claim-bearing decisions, companion standards, SGP evidence use, validation protocols, and public release claims. It does not add a public gate, does not alter the cascade, does not import an external ontology, and does not make empirical validation, legal certification, deployment certification, Tier 4, ProofPack, deterministic-selection, or automated moral-truth claims available.

AN.1 Purpose

Appendix AN protects RippleLogic runs from a methodological substitution failure: treating plausible output, coherent narrative, accumulated observations, governance approval, model tuning, institutional authority, or useful approximation as if it were sufficient support for a stronger claim than the evidence can bear.

AN.2 Required Methodological Integrity Record

When the Section 2.1F trigger holds, the PCC SHOULD include a MethodologicalIntegrityRecord and MUST include one for Tier 3 public or institutional claim-bearing runs and high-stakes Tier 2 runs where claim failure could change RF/NCRC, TRC, CSV, SGP, authority, deployment, public conformance, deterministic-selection wording, or validation status. The record contains: `claim_or_component`; `claim_type`; `claim_domain`; `warrant_domains`; `cross_domain_bridges`; `dependency_position`; `starting_assumptions`; `definition_or_operator_used`; `necessity_or_alternative_check`; `evidence_or_test_surface`; `falsification_or_revision_trigger`; `uncertainty_or_confidence_method`; `downstream_dependencies`; `re_derivation_scope_if_changed`; `reviewer_status`; and `required_claim_action`.

Dependency-localized falsification rule. Failure of a premise, primitive, operator, evidence source, control, or dependency SHALL reopen every materially dependent claim, score, gate result, authorization, and public statement. Components with an independently supported evidence path do not fail merely by association; they SHALL be evaluated on their own dependencies. This rule prevents both monolithic immunity and indiscriminate whole-framework rejection.

AN.3 Claim-Type Discrimination

Empirical claims require evidence contact. Causal claims require mechanism, counterfactual, intervention, or qualified causal warrant. Formal claims require explicit definitions, assumptions, and derivations. Operational claims require reproducible procedure and audit trail. Legal-regulatory claims require a declared jurisdiction, current authority source, applicability boundary, interpretation status, and contest or review path. Normative claims require declared ethical commitments, consistency, rights compatibility, stakeholder challenge, and auditability of application. Metaphysical horizon claims may guide humility and meaning, but they do not create Tier 1-3 evidence or gate authority.

Claim-domain and warrant-domain rule. A conclusion may draw on more than one warrant domain, but no warrant silently changes type. A legal rule may authorize, prohibit, allocate responsibility, or require testing without proving that a claimed technical capability exists. Scientific or engineering evidence may establish bounded capability or risk without creating legal authority or ethical legitimacy. Institutional approval, certification, compliance, commercial description, or anthropomorphic language cannot establish empirical truth, causal validity, physical safety, or reduced uncertainty by itself. A material cross-domain inference requires a declared bridge proposition, bridge evidence or rationale, uncertainty, and refusal or revision condition. BRIDGE_UNSUPPORTED blocks the stronger cross-domain claim; BRIDGE_PARTIAL permits only the supported narrower claim.

AN.4 Necessity and Alternative-Explanation Rule

If the run claims that a principle, definition, operator, threshold, indicator, configuration, or candidate result is necessary, the run must record why alternatives do not explain the same effect. If plausible alternatives remain open, the claim must be downgraded to assumption-bound, model-bound, contested, or exploratory status.

AN.5 Failure, Revision, and Re-Derivation Rule

A failed test, falsified prediction, contested evidence surface, failed dependency, or changed starting assumption cannot be hidden by silent model tuning. The record must declare whether the affected chain requires local revision, threshold revision, gate rerun, SGP rerun, CSV rerun, TRC reopening, RLS recalculation, public claim narrowing, release versioning, or refusal.

AN.6 Proportionality Rule

Methodological rigor is proportional to stakes. Ordinary Tier 1 and low-stakes Tier 2 uses may use lighter records. Tier 3, public conformance, institutional deployment, safety/reliability claims, SGP public claims, and high-consequence AI or physical/causal actions require stronger records. Proportionality may reduce documentation burden; it may not authorize a stronger claim than the evidence supports.

AN.7 Boundary Against Overclaim

This standard is not a claim that MathGov has discovered a final theory of reality, a universal physical stability law, or a complete empirical validation engine. It is a claim-discipline standard. Its function is to make assumptions, alternatives, tests, failures, dependencies, and required re-derivations visible before decisions are scored, published, deployed, or treated as governance-grade.

Appendix AO: Physical Execution Boundary (Normative Hardening Patch)

Status. Appendix AO is a v12.0 physical-execution boundary clarification. It is normative for consequence-bearing physical, biological, ecological, medical, infrastructure, cyber-physical, operational, or irreversible causal execution claims. It does not add a public gate, does not turn MathGov into a physics engine, formal-methods prover, simulator, medical device, legal authority, regulator, deployment-certification system, or domain-safety case, and does not make Tier 4 or ProofPack claims available.

AO.1 Boundary rule

Governance permission is not physical admissibility. Certification, documentation, organizational authority, procedural compliance, model conformity, approval, monitoring, or an ALLOW label MUST NOT be represented as proof that a physical execution can safely exist in physical reality.

For material consequence-bearing physical execution, MathGov may authorize only the governance status of a decision record unless a domain-appropriate physical or causal evidence surface also supports the execution claim. The framework may require, structure, audit, constrain, escalate, or refuse that evidence. It does not generate the physics, engineering, medical, biological, or safety proof by itself.

AO.2 Required routing

When physical or causal execution is material, the run MUST trigger the Physical/Causal Admissibility Evidence Profile inside Reality Grounding and CSV. If physical admissibility is not externally supported within the declared validity domain, the run MUST NOT claim physical safety. The permitted claim state is limited to GOVERNANCE_PERMISSION_ONLY, PHYSICAL_ADMISSIBILITY_NOT_ESTABLISHED, PHYSICAL_ADMISSIBILITY_CONTRAINED, or an equivalent narrower status until qualified domain evidence supports the stronger claim.

AO.3 ALLOW language

ALLOW means that the MathGov decision record did not find an upstream governance-gate bar under the declared evidence and scope. ALLOW does not mean that physical execution is safe, deployable, legally certified, medically cleared, or engineering-certified. A physical-execution claim requires the PC-AEP record and any external domain evidence named by the governing deployment context.

AO.4 Refusal condition

If the physical or causal evidence is absent, contradicted, outside validity domain, unrecoverably uncertain, irreversible without safeguards, unmonitored beyond tolerance, or dependent on unverified model assumptions, MathGov MUST narrow, control, redesign, escalate, or refuse the physical-execution claim. No RLS score may rescue the claim.

APPENDIX AP: CURRENT RULES-ONLY NAVIGATION (v13.0) - INFORMATIVE

This appendix is a navigation layer, not a second specification. It points implementers to the current controlling surfaces while keeping historical lineage separate.

Decision function	Current controlling surface	Minimum current rule
Ground claims	Sections 2.1A, 4, 5.7; PC-AEP v2.6; MFDI v2.6; Source-Coupling v2.6	Declare evidence, unknowns, validity domain, claim boundary, and refusal/revision trigger.
Protect rights	Sections 7 and Appendices B-C	Apply floor, categorical, and severe-hazard channels; use the same SevereRightsHazardActive predicate in activation and computation.
Bound ruin	Section 8 and Appendix D	Apply governed scenario discovery and CVaR; emergency-provisional action never becomes an ordinary pass.
Verify viability	Section 9; CSV v2.7	Test containment, dependencies, controls, reversibility, monitoring, source coupling, and execution conditions.
Rank survivors	Section 10; WDBIP v1.9	Populate the 7x7 welfare field without duplication, normalize by active

Decision function	Current controlling surface	Minimum current rule
		effective mass Q, and preserve uncertainty.
Resolve non-decisive outcomes	Sections 10.4 and 11	Refuse deterministic selection; record authority selection or redesign/escalation without calling judgment mathematical proof.
Authorize and execute	Sections 11-15; Agent v13.0; ripple.md v5.8	Selection is not authorization; authorization is not safe execution; tempo and responsibility continuity must remain reconstructable.
Replay and challenge	Reproducibility v1.7; run schema/validator; PCC	Pin evidence and parameters, preserve state transitions, report divergence, appeal, monitoring, and reopen triggers.

v13.0 carries forward the v12.7 integrity rules, including qualification continuity, assurance hardening, rights-effect non-dilution, stable effect identity and granularity, temporal-token integrity, typed likelihood, adverse-confidence review, distributional/subgroup closure, run-level mask symmetry, dependence-bounded decisiveness, finite robustness completeness, option-set closure, and normative-underdetermination reporting. These preserved rules retain the five-stage cascade and do not create a sixth gate.

APPENDIX AQ: CURRENT CANONICAL MACHINE-STATE AND MATURITY CONTRACT (v13.0) - NORMATIVE INTERFACE

This appendix carries forward the v12.4 machine-state and maturity contract as the current v13.0 interface. It is not historical fallback text.

The hash-pinned `CANONICAL_STATE_REGISTRY_v1.1.yaml` and `STATE_TRANSITION_MATRIX_v1.1.json` are the controlling machine-state surfaces for run-record v4.1. Documentation, schema, validator, workbook exports, and tests SHALL use the same emitted tokens. `CSV_REDESIGN` is accepted only as a deprecated migration alias and normalizes to `CSV_REDESIGN_REQUIRED` with a warning.

In-package token reference (Normative scope clarification). The external run-record v4.1 registry and transition matrix named above are not bundled or verified in this delivery. For the supplied scoped reference tests, Reproducibility and Use Standard Section 3 enumerates the in-package gate, decision and execution vocabulary; Canon Section 0.3, Section 11.5 and Appendix AF.2 control its meaning. The packaged state crosswalk records those sources and its subset explicitly. It is not a replacement for, or proof of conformance to, either absent external artifact. Any disagreement blocks the affected interchange claim until the controlling sources are reconciled.

The typed `MEASUREMENT_AND_PARAMETER_MATURITY_REGISTER_v1.0` separates empirical constructs, model components, normative priors, governance conventions, and unavailable design targets. A normative prior is not an unvalidated empirical measurement awaiting promotion. Framework-level validation or superiority claims cannot rest decisively on below-threshold empirical constructs or model components, and normative commitments cannot be laundered as empirical constants.

WDBIP owns WDBIP-specific profile, effect-token classification, token-level pathway, migration, schema, and validator surfaces. SGP owns MPS, FPP, GPR, SPR, ICP, RMCP, P100, and their misuse-review fields. The Canon owns Union Scope redundancy methods, allocation constraints, gate states, RLS weights, and the maturity-register type system.

APPENDIX AQ-1: CONFIGURATION-BOUND ASSURANCE AND READER NAVIGATION (CURRENT)

This release adds a cross-cutting assurance discipline without adding a gate, changing the seven Union Scopes, changing the seven Welfare Dimensions, changing numerical defaults, or altering the RG -> RF/NCRC -> TRC -> CSV -> RLS cascade.

Status and precedence. The configuration-binding, frozen-evaluation, requalification and Open Assurance Dependency requirements below are normative clarifications under Section 0.1. The abstract notes, process-model explanation and reader-navigation passages are informative. Neither part overrides a higher-precedence gate, equation or authority-separation rule.

Release-specific abstract notes relocated from the opening summary

The v13.0 / SGP v8.8 integration preserves proportional CSV qualification and the existing rights, ruin, and score hardening while replacing the legacy SGP protection surface interface with separate MPS, FPP, GPR, SPR, and ICP surfaces. It makes the protection-versus-authority separation computational rather than merely interpretive, requires interval sensitivity for SGP-aware residual welfare ranking, and prevents full protection from being misread as CMIU membership or consequential authority. It carries forward the v11.0 Term Discrimination and Semantic Stability Patch and makes the discriminator, dependency position, non-collapse risks, and refusal condition explicit for material terms so that governance does not proceed through collapsed language. The v13.0 integration preserves a claim-falsification and dependency-integrity discipline inside Reality Grounding and downstream validation: material claims must declare claim type, dependency position, assumptions, alternative explanations, evidence/test surface, falsification or revision trigger, and re-derivation scope when foundations change.

Configuration-Bound Assurance Principle

Every claim concerning capability, safety, physical admissibility, reliability, rights impact, containment, expected consequences, selectability, or authorization **MUST** be indexed to the declared system configuration, version, control stack, authority state, operating environment, and validity envelope under which the claim was established. A claim established for one configuration does not automatically transfer to a materially changed configuration.

Required distinction	Canonical rule
Configuration vs capability	A configuration enables and constrains a capability envelope; observing one successful process does not establish the complete envelope or its failure boundary.
Capability vs authority	Latent, available, enabled, authorized, selected, execution-approved, executed, and observed-outcome states are not interchangeable.
State vs transition vs post-state	A valid current state does not establish a valid transition; a valid transition does not establish a viable resulting state.
Qualification continuity	Qualification is configuration-bound, evidence-bound, authority-bound, control-bound, domain-bound, and time-bound. Material change suspends inherited authority until requalification.
Structural control preference	Where feasible and rights-compatible, remove or bound the enabling pathway before relying only on behavioral expectation, confidence scoring, or downstream detection.
Future option space	Tier 2-3 CSV analysis records whether an option expands, preserves, narrows, or irreversibly collapses future rights-compatible and structurally viable choices.

Recursive configuration-process model

For assurance purposes, MathGov uses the bounded recursive relation:

$$(C_t, E_t) \rightarrow K_t \rightarrow A_t \rightarrow P_t \rightarrow (C_{t+1}, E_{t+1})$$

where C_t is the declared configuration, E_t the relevant environment, K_t the capability envelope, A_t the admissible and authorized action set, P_t the selected and executed process, and C_{t+1} the resulting configuration. Current configuration constrains current capabilities; executed processes can alter future configuration and future capabilities.

Frozen evaluation binding

Before consequential comparison, the decision question, configuration, system boundary, evidence schema, stakeholder scope, rights interpretation, scenario families, materiality thresholds, uncertainty treatment, control requirements, weighting protocol, stopping conditions, and authorization pathway **MUST** be frozen and versioned for that run. A material post-result change creates a new run or a declared revision branch; it may not be represented as uninterrupted continuation of the original evaluation.

Requalification triggers

Requalification is required when a material model, software, hardware, sensor, tool, permission, control, environmental, legal, authority, stakeholder, evidence, dependency, risk, monitoring, shutdown, or operating-envelope condition changes; after a material incident or near miss; or when cumulative wear, drift, capability emergence, or new evidence weakens a prior warrant.

Open Assurance Dependency Rule

A proprietary, opaque, non-reproducible, or independently uninspectable component may contribute evidence, but it may not serve as the sole warrant for a Rights Floor, catastrophic-risk, structural-viability, selectability, or authorization conclusion. Decisive reliance requires a declared version and configuration, inputs and outputs, validity domain, test conditions, known failure modes, uncertainty bounds, challenge interface, provenance, fallback behavior, and independent evaluation proportionate to stakes.

Implementation-milestone claim boundary

Successful compilation, synthesis, simulation, hardware realization, integration, or benchmark execution establishes only the milestone actually demonstrated. It does not by itself establish empirical superiority, domain validity, physical safety, ontological truth, rights compatibility, lawful authority, or deployment readiness. The scientific sequence is: derive, implement, attempt to falsify, measure, publish, independently replay, and revise.

How to read this Canon

Read the current rules first. The architecture and active obligations are controlling; release-history material is retained for provenance and migration only. This DOCX is the delivered prose/table master under Appendix RELEASE. Repository and website publication require separate synchronized upload and retrieval checks, not a change of source authority by implication.

Current rules-only navigation

This navigation layer is not a second specification. It points implementers to the current controlling surfaces while keeping historical lineage separate.

Decision function	Current controlling surface	Minimum current rule
Ground claims	Sections 2.1A, 4, 5.7; PC-AEP v2.6; MFDI v2.6; Source-Coupling v2.6	Declare evidence, unknowns, validity domain, claim boundary, and refusal/revision trigger.
Protect rights	Sections 7 and Appendices B-C	Apply floor, categorical, and severe-hazard channels; use the same

Decision function	Current controlling surface	Minimum current rule
		SevereRightsHazardActive predicate in activation and computation.
Bound ruin	Section 8 and Appendix D	Apply governed scenario discovery and CVaR; emergency-provisional action never becomes an ordinary pass.
Verify viability	Section 9; CSV v2.7	Test containment, dependencies, controls, reversibility, monitoring, source coupling, and execution conditions.
Rank survivors	Section 10; WDBIP v1.9	Populate the 7x7 welfare field without duplication, normalize by active effective mass Q, and preserve uncertainty.
Resolve non-decisive outcomes	Sections 10.4 and 11	Refuse deterministic selection; record authority selection or redesign/escalation without calling judgment mathematical proof.
Authorize and execute	Sections 11-15; Agent v13.0; ripple.md v5.8	Selection is not authorization; authorization is not safe execution; tempo and responsibility continuity must remain reconstructable.
Replay and challenge	Reproducibility v1.7; run schema/validator; PCC	Pin evidence and parameters, preserve state transitions, report divergence, appeal, monitoring, and reopen triggers.

For the full navigation note and current machine-state interface, see Appendices AP-AQ.

Maturity and assurance ladders: type crosswalk (Informative)

These ladders measure different properties and MUST NOT be collapsed or translated automatically. A higher state on one ladder does not substitute for a required state on another. Tier-specific minima remain controlled by the cited source.

Family	Ladders	What they measure	Non-substitution rule
Decision depth	Tier 1-4	Required process depth, evidence, participation, and authority proportional to stakes.	Tier does not establish replay, measurement, implementation, or validation maturity.
Record and replay	R0-R4; L0-L3; V0-V2	Replay maturity, wrapper/conformance completeness, and validator capability.	A complete record or passing validator does not establish evidence truth or ethical correctness.
Method and measurement	SCI-0-4; UCI-M0-M3; K0-K4; BEM-0-4	Source coupling, UCI measurement, kernel quality, and biological-evaluation maturity.	One mature method surface cannot compensate for an immature or missing required surface.
Sentience and role evidence	MPS-NE/0-4; GPR 0-5; SPR 0-5; RMCP 0-4	Possible moral patienthood, governance participation readiness, stewardship readiness, and reality-management capacity.	Protection is not authority; capability is not authority; missing evidence is not zero.

Document map

Reading purpose	Primary location
Problem, ontology, and normative commitments	Sections 1-3
Architecture, states, and end-to-end procedure	Section 4
Welfare construction and propagation	Sections 5-6
Rights, catastrophic risk, and viability gates	Sections 7-9
Residual ranking and non-decisive handling	Sections 10-11

Reading purpose	Primary location
Records, audit, governance, validation, and limitations	Sections 12-17
Detailed formulas, interfaces, examples, migration, and conformance material	Appendices A-AR

Compact implementation and reproducibility contract

The controlling reading specification is

Core_15/MATHGOV_REPRODUCIBILITY_AND_USE_STANDARD_v1.7.docx. The full Canon controls the compact operating contract. The supplied MATHGOV_REPRODUCIBILITY_AND_USE_STANDARD_v1.7.docx is the current local operating companion; docs/reproducibility/MATHGOV_REPRODUCIBILITY_AND_USE_STANDARD_v1.5.md is an external prior source requiring synchronization before promotion. Identify the run and authority boundary; lock decision-material evidence and parameters; complete RG, RF/NCRC, TRC and CSV in order; form the selectable set; apply RLS only to survivors; distinguish decisive selection, authority selection, refusal, redesign, delay and emergency-provisional action; separate selection from execution authority; record controls and monitoring; and reopen affected stages after material change.

The companion docs/implementation/NORMATIVE_KERNEL_INDEX_v1.1.yaml maps each compact rule to the controlling Canon source. schemas/mathgov_run_record_v4_1.schema.json and release/VALIDATE_MATHGOV_RUN.py mechanically check record completeness and cascade contradictions. Passing those checks establishes only record/cascade conformance for the supplied inputs. It does not verify evidence truth, measurement validity, legal authority, physical safety, deployment readiness, or moral truth.

APPENDIX AR: VERSION HISTORY AND MIGRATION NOTES (v13.0) - INFORMATIVE

This appendix preserves detailed release notes for provenance, audit, and migration. It does not override the current controlling rules. Package-level history is also maintained in CHANGELOG.md and LINEAGE.md.

AR.0 Release Package, Source Hierarchy, and Publication State

Package line: MathGov/RippleLogic v13.0, Core 15, Canon v13.0 / SGP v8.8. The edition is distinct from the release-cycle label and the manifest-identified build.

Release-cycle date convention.

The 2026.09 label is a planned release-cycle identifier, not the month in which every source file was authored. Same-version patches made after the Canon date MUST be recorded in the release ledger, trigger dependency-localized requalification, and regenerate all manifests and hashes before publication.

Package-note: v13.0 preserves the five-stage cascade, seven Union Scopes, seven Welfare Dimensions, rights floors, and SGP protection invariants. It carries forward action-bound qualification continuity, executable current-state/transition/post-state checks, outcome observation, and dependency-localized requalification while retaining configuration-bound assurance, capability-state separation, frozen evaluation binding, structural-control preference, future safe option-space analysis, and the Open Assurance Dependency Rule. It pins SGP v8.8, WDBIP v1.9, ripple.md v5.8, Agent System v13.0, Reproducibility Standard v1.7, Cascade v2.9, CSV v2.7, RLS Validation v2.9, Primer v4.7, Public Introduction v13.0, PC-AEP/MFDI/Source-Coupling v2.6, and Aligners Sheet v5.9. No sixth gate, eighth welfare dimension, authority grant, empirical-validity claim, legal authorization, physical-safety certification, consciousness detector, ProofPack claim, or Tier-4 claim is created.

Source-hierarchy note. MathGov remains the umbrella framework and RippleLogic its decision architecture. For this delivered package, the Core_15 DOCX files are the controlling prose and table masters under Appendix RELEASE. Generated Markdown and PDF files are reading projections; discrepancies require repair from these masters. Internal normative precedence remains Appendix AC / Section 0. A repository may publish a synchronized source tree under its declared release policy, but absent repository files or old public mirrors do not override this package, and no unperformed promotion or machine-source parity is asserted.

Publication-state boundary. A source package may be internally release-ready before repository tags, public websites, or mirrors are synchronized. Public statements that this is the currently published canonical release MUST be checked against the accessible release tag, manifest, and SHA-256 set. Site publication is an external release action, not evidence created by this document.

Current release at a glance

Item	Current release position
Governing version	RippleLogic Canon v13.0 within MathGov Core Release 2026.09 / SGP v8.8
Governing architecture	RG -> RF/NCRC -> TRC -> CSV -> RLS
Preserved rules and candidate changes	RightsEffectToken non-dilution; stable effect identity and token granularity; temporal integrity; typed likelihood and probability ownership; adverse-confidence counterfactual review; mandatory distributional and subgroup routing; comparison-mask symmetry; reference and baseline families; dependence- and numerical-guard reporting; robustness completeness; weight-profile and normative underdetermination handling; option-set closure; and corresponding regression vectors. Candidate corrections add every-contender discrimination, scoped aggregate-risk checks, explicit metric contracts and the separately documented companion repairs without removing these preserved safeguards.
Preserved invariants	Five stages, seven Union Scopes, seven Welfare Dimensions, non-compensatory rights, catastrophe non-dilution, SGP protection/authority separation, refusal under underdetermination
Claim boundary	Current Tier 1-3 research and teaching specification, with a frozen worked example and scoped reference tests. Public upload is a separate promotion action. Unresolved empirical, native-engine and external-interface evidence remains explicitly disclosed; no high-stakes autonomous deployment approval is claimed.
Full version history	Appendix AR, CHANGELOG.md, and LINEAGE.md

AR.7 RippleLogic Canon v12.7 Release (Informative Migration Record) (PRIOR RELEASE)

Release date: 2026-08-26. Canon v12.7 is a bounded integrity-hardening release. It preserves RG → RF/NCRC → TRC → CSV → RLS, the 7×7 field, RLS as RippleLogic Score, HDW floors, PLSS, SGP type separation, profile-first reporting, non-compensatory rights, separate ruin control, and framework/authority/execution separation.

The release adds: RightsEffectToken non-dilution; stable effect identity and token granularity; temporal token integrity; typed likelihood and probability ownership; adverse-confidence counterfactual; mandatory distributional/subgroup routing; comparison-mask symmetry; reference and baseline families; dependence and numerical-guard reporting; finite robustness completeness; weight-profile status and normative underdetermination; option-set closure; and formal regression vectors.

The release does not claim empirical validation, calibrated universal welfare measurement, legal authority, physical safety, deployment readiness, Tier 4, ProofPack completion, or moral truth. v12.6 calculations remain historical lineage artifacts and MUST NOT be silently reinterpreted as v12.7 runs.

AR.7A Same-Version Final Improved Build (2026-09-03) (PRIOR RELEASE)

Build status. This same-version correction and scientific-clarity build retains RippleLogic Canon v12.7 and all companion semantic versions. It corrects stale current-facing pins, separates FPP/full inclusion from MPS evidence, adds explicit instrumental tool-role and bidirectional artificial-system self-report boundaries, tightens functional-understanding claims, and synchronizes the derived public and workbook interfaces.

No gate, threshold, equation, Union Scope, Welfare Dimension, rights floor, TRC corridor, CSV status, RLS formula, SGP protection floor, or authority rule is changed. The build becomes a published canonical artifact only after its exact governing sources and reading mirrors, release manifest, checksum ledger, and publication surfaces are regenerated and verified.

AR.0A v12.2-v12.4.1 Consolidated lineage bridge

The v12.2 line consolidated source integrity, SGP v8 integration, operational assurance, formula interfaces, and workbook integrity. The v12.4 and v12.4.1 lines synchronized executable-state, consequence-tempo, responsibility-continuity, WDBIP, release-manifest, and workbook surfaces. Their detailed reports remain in archive/release_history and archive/release_v12_4_1_legacy. These lineage entries do not govern current v13.0 interpretation.

AR.1 v12.4 Calculability, Type-Integrity, and Synchronization Release

v12.4 is a coordinated minor release because it changes a load-bearing rights-risk predicate, corrects a reference uncertainty derivation, completes the SGP RMCP misuse-record interface, and synchronizes every active core surface. The governing architecture remains $RG \rightarrow RF/NCRC \rightarrow TRC \rightarrow CSV \rightarrow RLS$. The severe-rights-hazard channel now uses one activation predicate in both review and computation; a hazard cannot activate under Section 7.4.3 and then disappear from v_r^{risk} . Method B reference-vector reconstruction now divides by the declared confidence factor when recovering pre-confidence contribution mass. Registered impact-stream notation uses $h_{\{k,base\}}$ and $h_{\{k,welfare\}}$; legacy s_k and m_k aliases are historical only.

The release also adds a current-rules-only navigation appendix and machine-checkable agreement tests for severe-rights activation and SGP RMCP required fields. These are integrity improvements, not evidence that the governing thresholds or instruments are empirically calibrated.

AR.2 v12.4.1 Release and Workbook Integrity Patch (Lineage)

v12.4.1 is a bounded synchronization and release-engineering patch. It repairs exact companion pins, current-package filenames, public release labels, workbook formula and named-range integrity, CVaR provenance, sanity-check completeness, and the 15-file review bundle. It does not change the canonical cascade, rights floors, TRC mathematics, CSV semantics, RLS equation, Union Scopes, Welfare Dimensions, or SGP protection architecture.

AR.3 v12.4 Integration and Executable-Integrity Note (Lineage)

v12.4 preserves the governing cascade $RG \rightarrow RF/NCRC \rightarrow TRC \rightarrow CSV \rightarrow RLS$, the seven Union Scopes, the seven Welfare Dimensions, the Rights Floor, TRC, CSV, RLS equations, and all SGP protection floors. It integrates four bounded improvements: (1) WDBIP v1.4 as a Canon-subordinate welfare-measurement companion beneath RLS; (2) consequence-tempo and responsibility-continuity assurance through ripple.md v5.2, CSV v2.1, and Agent System v12.1; (3) a single versioned state registry and deterministic stage-transition algebra; and (4) typed measurement and parameter maturity disclosure. These additions do not prove empirical validity, legal authority, physical safety, or moral truth.

AR.4 v12.2.1 Retained Patch Integrity Note (Lineage)

This patch does not change the five-stage architecture. It repairs formula and interface integrity: stream notation now uses $h_{\{k,x\}}$ consistently; MPS intervals route to governed inclusion hypotheses rather than numeric endpoint multipliers; U7 Equity distinguishes NA from UNKNOWN; RLS companion formulas retain normalization by Q; CSV tier obligations and status tokens are synchronized; current worked examples disclose all three NCRC channels; and workbook sanity, publishability, and PLSS-limit records are mechanically reconciled.

APPENDIX AX: NORMATIVE AND CLAIM-AUTHORITY CLARIFICATION PATCH

This appendix is normative for interpretation and preserves the existing five-stage cascade, equations, state semantics, Union Scopes, Welfare Dimensions, thresholds, and companion interfaces. It adds no sixth gate and changes no existing numerical default.

AX.1 Operational relational premise

Within a declared decision domain, scope, horizon, relation family, materiality threshold, and evidence sensitivity, a RippleLogic run MUST identify materially affected beings, enabling systems, dependencies, consequence pathways, representation duties, stewardship obligations, and foreseeable risks. Nondetection is not automatically absence; uncertainty is not permission to invent relations.

This operational premise does not establish universal relationality, unrestricted causal reachability, ontological monism, one substance, one consciousness, physical infinity, or AIU-O. Relational Union and AIU-O remain optional foundational research questions and are not prerequisites for running or adopting MathGov.

AX.2 Explicit normative bridge register

The Minimal Normative Axiom remains the compact public entry point. For constitutional interpretation, the following commitments and bridge premises state the fuller reasons for qualification-before-ranking. They justify a family of rights-constrained, ruin-aware, relationally attentive governance architectures; they do not uniquely entail MathGov's exact taxonomy, gate order, thresholds, weights, scoring formula, or institutional implementation.

ID	Commitment or bridge premise	MathGov consequence
N1	Moral relevance: the welfare, rights, freedom, continuity, agency, and enabling conditions of materially affected beings are morally relevant.	Map affected beings and enabling conditions.
N2	Non-compensability: some violations of life, bodily integrity, liberty, dignity, fair process, information rights, and essential ecological conditions cannot be rescued merely by aggregate benefit.	Maintain a non-compensatory Rights Floor.
N3	Differentiated protection and non-domination: protection responds to credible evidence of sentience, vulnerability, dependency, agency, consent capacity, and exposure; difference does not justify arbitrary domination.	Use SGP and protection typing without converting protection into authority.
N4	Ruin asymmetry: catastrophic, irreversible, correlated, and existential harms warrant stronger precaution than ordinary reversible loss.	Apply non-linear tail-risk control before residual ranking.
N5	Conflict and emergency handling: justified limitations must be distinguished from prohibited violations and must be lawful, necessary, proportionate, temporary, monitored, reviewable, and reversible where feasible.	Govern emergencies explicitly; do not normalize exceptions.
B1	Material affectability: operational consideration requires a credible pathway of effect, dependency, representation, stewardship, or foreseeable risk.	Include by evidenced pathway, not mere co-membership.

ID	Commitment or bridge premise	MathGov consequence
B2	Agency under uncertainty: responsibility scales with power, knowledge, role, and capacity to foresee, control, or mitigate effects.	Record responsibility, authority, and control asymmetries.
B3	Public justification and correction: coercive or high-impact decisions require contestable reasons, auditable evidence, legitimate representation, remedy, and revision.	Require challenge, appeal, redress, and correction pathways.

AX.3 Protection-type separation

Protection type	Primary basis	Authority implication
Moral-patient protection	Credible sentence, valence, suffering, welfare, or vulnerability evidence	No automatic authority
Rights protection	Constitutional, legal, or declared moral status	No automatic authority
Agency protection	Autonomy, consent, freedom, and capacity to pursue ends	May support participatory standing, subject to lawful role
Ecological protection	Life-support function, integrity, biodiversity, and stewardship	No authority grant
Structural protection	Institutional continuity, critical infrastructure, or systemic viability	No authority grant
Representational protection	Material exposure combined with inability to participate directly	Requires legitimate representation or advocacy
Decision authority	Lawful mandate, competence, role, and accountability	Bounded authority only

No record may collapse moral-patienthood, rights-bearing status, agency, participation, protection, and decision authority into one scalar category.

AX.4 Unioning and residual-harm rule

Unioning is the disciplined search for cooperative, Pareto-improving, harm-reducing, burden-sharing, and structurally safer redesigns before accepting residual trade-offs. It does not assume that every viable decision can make every Union positively affected. Residual non-rights harms MUST remain visible, bounded, justified, minimized, monitored, and, where appropriate, remedied. Rights failures and prohibited catastrophe exposure remain non-compensatory and cannot be redescribed as acceptable residual harm.

AX.5 Profile-first reporting

For Tier 2 and Tier 3 runs, the disclosed 7 x 7 profile, subgroup and worst-affected-group slices, uncertainty, evidence maturity, non-maskable findings, sensitivity results, and incomparability or non-decisiveness status MUST be presented before any scalar RLS summary. RLS is an authorized summary of a disclosed profile, not a substitute for that profile.

FINAL DOCUMENT COMPLETION STATEMENT

RippleLogic v13.0 is supplied as the current Tier 1-3 research and teaching specification. It preserves RG -> RF/NCRC -> TRC -> CSV -> RLS, seven operational scopes, the SGP binding interface, PLSS/FLGG, reference-structure discipline, refusal logic and audit requirements. The current adjudication and verification reports identify corrections, executed checks and remaining evidence limits. This delivery is prepared for review, reading and scoped testing; live publication is not asserted.

Appendix AX is the final normative clarification appendix. This completion statement is followed by Appendix RELEASE, which records identity, source authority and provenance; the document therefore does not end at

AX. Labels AS-AW are intentionally unused. The manifest and SHA-256 ledger identify this build. The supplied reference code implements only its declared subset: no complete external-registry conformance, empirical validation, Tier 4 or ProofPack readiness, deployment or legal certification, universal UCI measurement, completed biological SGP measurement, or automated moral truth is claimed.

APPENDIX RELEASE: Identity, Source Authority and Revision Record

Framework release: MathGov/RippleLogic v13.0. Component: RippleLogic Canon v13.0. Edition-origin preparation: 10 September 2026; the separately identified correction build is dated below. Two-part major.minor component versions are used; preserved historic identifiers are not renumbered.

Status: integrated research and teaching specification with a bounded worked-example and scoped reference implementation. Readiness is limited to the checks in Reports/Verification_and_Readiness.md. No empirical validation, independent human validation, full production runtime, physical-safety certification, legal authority, deployment authorization or Tier-4 ProofPack status is implied by the edition number.

Delivery identity: MathGov/RippleLogic v13.0, build MG-RL-13.0-20260923-RELEASE-G. Prepared 23 September 2026 from the supplied frozen publication archive. Only reproduced defects and approved clarifications were patched; component editions are unchanged. The manifest identifies the exact current bytes and the preserved baseline. This is a new correction build, not a silent replacement of the earlier frozen artifact. Live publication is not asserted.

Component/build identity. Aligners Sheet v5.9 retains its component edition but carries this correction build's identifier. Its numerical inputs, calculation formulas and worked verdict are unchanged; label, runtime-token and integrity-snapshot corrections are itemized in Verification/Exact_Workbook_Changes.json. Exact hashes, rather than filenames or edition labels alone, distinguish the new bytes from Release D and the previous frozen publication.

Verification boundary: Reports/Verification_and_Readiness.md and Verification/Final_G/ contain the current build's executed checks and limitations. Earlier verification and the rejected compatibility experiment belong to the preserved baseline and do not certify changed bytes. Cache-independent replay, stored-formula and native-engine results are separate evidence surfaces. Microsoft Excel parity, full external-registry conformance, empirical validation and production authorization are not asserted.

Package source authority: Core_15 contains 14 DOCX specification masters and one XLSX Aligners frozen worked-run workbook. DOCX files control prose, equations and tables subject to Canon and SGP ownership. The workbook controls only its disclosed exemplar surfaces. Sources, Reading_HTML and Reading_PDFs are generated reading projections, not competing normative masters. The manifest and hashes identify exact bytes; supplemental code and schemas govern only their documented subset. Conflicts require recorded correction against the controlling source.

License and reuse: consult the package-level LICENSE and NOTICE and retain applicable component-specific and third-party notices. This pointer does not override a valid exception or grant rights over separately supplied private or third-party review material.

External repository publication and any absent legacy schema/validator/registry are separate authorities/evidence surfaces. The local package does not claim to update a remote repository or reproduce unavailable implementations. Current source pointers in this appendix replace prior front-matter edition pointers for this package; accurate historical source references below are retained as lineage only.

Preserved invariants: RG -> RF/NCRC -> TRC -> CSV -> RLS; rights non-compensation and unallocated rights effects; existing seven rows and seven dimensions; existing SGP MPS/FPP/GPR/SPR/ICP/RMCP separation; probability ownership, uncertainty, stability and every-contender rules; selection, authority and execution separation.

Component	Current edition
RippleLogic Canon	v13.0
Sentience Gradient Protocol	v8.8

Component	Current edition
ripple.md Standard	v5.8
RippleLogic Agent System	v13.0
CSV Gate Standard	v2.7
RippleLogic Cascade Standard	v2.9
MathGov Reproducibility and Use Standard	v1.7
Welfare Dimension Boundary and Interaction Protocol	v1.9
RLS Validation Protocol	v2.9
RippleLogic Foundations Primer	v4.7
MathGov Public Introduction	v13.0
Physical/Causal Admissibility Evidence Profile	v2.6
Methodological Falsifiability and Dependency Integrity Standard	v2.6
Source-Coupling Integrity Standard	v2.6
RippleLogic Aligners Sheet	v5.9

Change locations and rationales: Reports/Audit_Adjudication_and_Changes.md; exact edits: Verification/Exact_Document_Changes.json and Exact_Workbook_Changes.json. Navigation and metadata records are separate. Prior releases and feedback are preserved in the complete provenance archive.

Current proposals are not universal truths. Taxonomy maximality, continuous-time propagation, cross-substrate cardinal welfare, generic susceptibility/shield formulas, autonomous recovery, vendor infrastructure, zero-knowledge circuits and hardware meshes remain unvalidated unless independently demonstrated under a scoped implementation profile.

Preserved Baseline Release Material (Historical; Non-Controlling)

The following blocks are relocated intact from the recovered baseline. Their edition numbers, release-readiness wording and external-source limitations describe that historical candidate, not current component identity or newly executed verification. Governing current metadata is the matrix above.

Release control	Historical value
Component	RippleLogic Canon v12.8-rc2
Release	HISTORICAL (NON-CONTROLLING): MathGov Core Release 2026.09 — RippleLogic Canon v12.8-rc2 / SGP v8.7-rc2
Architecture	RG → RF/NCRC → TRC → CSV → RLS
Role	Principal normative decision architecture. It governs cascade semantics, equations, thresholds, source hierarchy, records, audit flags, and claim boundaries.
Claim boundary	Controlled Tier 1–3 research/specification candidate; not empirical validation, legal authority, physical-safety

Release control	Historical value
	certification, deployment authorization, Tier 4, or moral truth.
Source/render parity	Candidate files, internal identities, pins, manifest and hashes must agree. Repository semantic-source and publication-mirror synchronization is required before promotion; parity with an absent candidate Markdown source is not claimed.

Historical Release Integration (v12.8-rc2)

This review candidate advances RippleLogic Canon from the preserved audited v12.7 baseline through v12.8-rc1 to v12.8-rc2, with SGP v8.7-rc2 and the exact companion matrix in the candidate manifest. It retains the rc1 rights non-dilution, aggregate-risk, every-contender, validation, runtime-interface, and worked-aid corrections, and adds bounded rc2 clarifications for propagation scope, conformance vectors, command-envelope determinacy, RMCI_L zero-bound handling, workbook integrity, navigation, and release verification. Historical releases remain unchanged and interpretable. This candidate is not a published release; empirical validation, production implementation, external repository parity, and deployment authorization remain separately evidenced claims.