

# MathGov and RippleLogic Foundations Primer

| Integrity surface  | Current requirement                                                                                                     | Claim boundary                                                                        |
|--------------------|-------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------|
| Framework identity | MathGov is the umbrella ethical-governance framework; RippleLogic is its decision architecture and method.              | The Primer is informative and yields to the Canon.                                    |
| Public doorway     | Reality -> Rights -> Ripples and Qualify -> Rank provide the shortest safe explanation.                                 | Teaching simplification may not erase TRC, CSV, uncertainty, or authority boundaries. |
| Ontology boundary  | Union-Based Ethics is a declared normative foundation; AIU remains a philosophical horizon rather than empirical proof. | No metaphysical claim substitutes for evidence.                                       |

## What a system can do, may do, and did do

MathGov separates three questions that are often blurred:

1. **What can the present configuration do?** This is a capability question.
2. **What may it do?** This is an admissibility, rights, risk, viability, and authority question.
3. **What did it actually do, and what changed afterward?** This is an execution and observed-consequence question.

A tested capability belongs to the configuration in which it was demonstrated. Changes to models, tools, permissions, sensors, controls, institutions, law, environment, or evidence can change the answer and require requalification. This is why MathGov treats alignment as a continuing relationship with reality rather than a one-time certificate.

## AI systems: tool-role, capability, welfare, and authority

An instrumental tool-role describes how a system is deployed: one or more actors use it to extend their capacities toward assigned ends within an arrangement of objectives, permissions, responsibility, and control. Tool-role does not tell us how intelligent or agentic the configured system is, whether it has welfare-bearing experience, how it should be protected, or what legitimate authority it may hold.

AI governance therefore has two connected but non-collapsing tracks. Capability–Consequence Governance asks what the exact system configuration can do, what consequences it can cause, and what evidence, controls, permissions, accountability, and authority are required. Possible-Welfare Governance asks how the system should be treated if evidence supports—or cannot safely exclude—welfare-bearing states. Capability does not establish sentience. Possible sentience does not establish competence or authority. These are explanatory tracks inside the existing MathGov architecture, not new gates.

## Configuration-Bound AI Characterization Summary (Informative)

For communication and review, the following read-only summary may be assembled from existing records. It is not a new status object, scalar, gate, or independent authority source.

| Surface               | Governing question                                                                  | Existing owner or record                      |
|-----------------------|-------------------------------------------------------------------------------------|-----------------------------------------------|
| Exact system identity | Which model, configuration, environment, permissions, and time are being evaluated? | ConfigurationRecord; Reproducibility Standard |
| Instrumental role     | How and by whom is the system deployed                                              | Canon Category Grounding; Agent               |

|                          |                                                                                        |                                                        |
|--------------------------|----------------------------------------------------------------------------------------|--------------------------------------------------------|
|                          | and controlled?                                                                        | Computational Context Preflight                        |
| Functional capability    | What has the exact configuration demonstrated?                                         | CapabilityAuthorityDecompositionRecord; ICP where used |
| Functional understanding | What target, relation, and domain are supported under what tests?                      | ICP; MFDI capability-claim record                      |
| Agency and autonomy      | How does it select, persist, revise, and act within an environment?                    | Agent System; capability and role records              |
| Phenomenal evidence      | What evidence supports welfare-bearing experience, and what remains unresolved?        | MPS or MPS-NE / NOT_EVALUATED                          |
| Protection               | What safeguards follow from evidence, vulnerability, stakes, and uncertainty?          | MPS protection posture; FPP                            |
| Participation            | How can the entity participate directly, with support, or through representation?      | GPR                                                    |
| Consequential role       | Is it eligible for one specifically named role under constraints?                      | SPR                                                    |
| Lawful authority         | Who may legitimately authorize the decision or action?                                 | External lawful mandate; AuthoritySelectionRecord      |
| Runtime state            | Is capability latent, available, enabled, authorized, selected, approved, or executed? | Agent System capability-state record                   |
| Observed outcome         | What actually happened, and what must be reopened or repaired?                         | OutcomeObservationRecord; NCAR reflection              |

The summary must reference the governing records and hashes where available. No field may borrow authority from another: capability is not permission, protection is not command power, and execution is not an acceptable outcome by itself.

MathGov Institute for Ethical Systems Design

## Consistency, Tier Integrity, Rights, Ruin, Tempo, and Responsibility Note

MathGov now explicitly separates governance approval from physical or causal admissibility. When a material physical or causal action is proposed, the run should record the Physical/Causal Admissibility Evidence Profile inside Reality Grounding and CSV. The profile asks: what model or causal basis is used, where it is valid, what boundary conditions hold, what uncertainty remains, what can fail, what harm is reversible or irreversible, what verification or expert warrant exists, how the action will be monitored or stopped, what remains unknown, and whether the claim action is proceed, narrow, control, redesign, escalate, or refuse.

This improves MathGov without pretending that MathGov itself is a physics engine, safety certification, medical review, legal approval, or formal proof system.

## ABSTRACT

This primer is the human doorway into MathGov, RippleLogic, and the Sentience Gradient Protocol. It explains why the framework exists, what problem it addresses, what it claims, what it does not claim, how the main pieces fit together, and why material terms must be discriminated before they can govern decisions. It is informative only: the Canon and companion specifications remain the source of normative force.

## Plain-language summary

The Rights Floor (RF) is the public-facing name for the Non-Compensatory Rights Constraint (NCRC). In plain terms, the Rights Floor means that protected rights violations cannot be averaged away by downstream benefits, majority preference, or an attractive RippleLogic Score.

Decisions ripple outward. They affect people, institutions, ecosystems, future generations, and sometimes intelligences that are not yet well understood. The MathGov framework, using RippleLogic, gives readers a structured five-level method: Reality Grounding, Rights Floor (RF), TRC, CSV, and RLS. In simpler words: ground the facts, protect the Rights Floor, bound ruin, contain and verify the whole, then choose the strongest conditional residual ripple among selectable options, if decisive. The framework makes ethical reasoning auditable rather than merely rhetorical. This candidate integrates the current release discipline while preserving a simple public discipline: ground the words before trusting the score.

## How to read this primer

Begin with Sections 1 and 2 for the problem and basic identity of the framework. Then read Sections 3, 3A, 3B, 3C, and 22 for layer discipline, primitive audit, reference-structure reasoning, Category Grounding, and Term Discrimination. Use later sections to understand the cascade, SGP, audit posture, and release boundaries.

## Section 1. The Problem MathGov Is Built To Address

Human authority alone is not stable enough for civilization-scale decision making. Humans can be wise, compassionate, courageous, and visionary. Humans can also be distorted by fear, greed, coercion, propaganda, concentrated power, institutional drift, tribal loyalty, short-term incentives, and survival pressure.

That is the surface problem. The deeper problem is structural: complex systems frequently reward local optimization while concealing larger-system harm.

A company can improve quarterly results while burning out its people. A nation can grow GDP while degrading the ecosystems supporting its survival. A government can claim stability while suppressing rights. An AI lab can increase capability while making alignment harder to reason about. From inside each system, the move can look rational. Zoom out far enough, and the pattern becomes dangerous.

Local optimization. Global degradation.

That is one of the central fractures MathGov and RippleLogic are designed to address. The goal is not to replace human judgment with a machine. The goal is to make judgment more transparent, auditable, less corruptible, less arbitrary, less captured by power, and more accountable to the full field of consequences.

## Section 2. What MathGov Is, and Is Not

### 2.1 What MathGov is

MathGov is a reality-seeking decision architecture. RippleLogic is its decision engine. SGP is its moral status and protection protocol. Together they constitute a proposed ethical operating-system architecture for governance, AI alignment, institutional design, and high-stakes decision-making under complexity.

### 2.2 What MathGov is not

MathGov is not a claim of final truth. It is not a substitute for human judgment, lawful authority, or democratic legitimacy. It does not claim that all values reduce perfectly to numbers. It does not claim that current AI systems are conscious. It does not claim to have solved alignment. It does not claim universal consent across cultures.

MathGov proposes the design hypothesis that ethical and governance decisions can become harder to corrupt, harder to hide, harder to capture, and easier to test against reality when their primitives are declared, their assumptions are exposed, their weights are visible, their rights floors are protected, their uncertainty is modeled, and their reasoning is auditable.

Architecture-final is not ProofPack-final. A framework can be coherent, specified, and useful before it is empirically validated, schema-complete, validator-complete, or Tier-4-ready; those stronger statuses require separate evidence and replay materials.

## 2.3 Layer discipline

MathGov becomes clearer when each kind of statement is kept in its proper layer. It does not treat ontology, ethics, mathematics, evidence, governance, and interpretation as the same thing.

Operational ontology. Union-Based Reality describes the decision-relevant structure of interdependence: beings and systems exist inside relationships, dependencies, and consequences. This is an operational claim for governance-relevant domains, not a final metaphysical proof of all being.

Reality, decision-relevant reality, reality surface, evidence, and interpretation are related but not identical. Reality is the evaluator-independent structure of constraint and consequence that holds whether or not any agent perceives, admits, models, or has power over it. Decision-relevant reality is the portion of that structure that materially bears on the decision. A reality surface is the portion actually captured, represented, or modeled in the run. Truth, for a decision claim, means correspondence to decision-relevant reality as supported by the declared reality surface and evidence trace. Evidence is the trace by which a claim is tested against the relevant surface. Interpretation is the judgment that connects evidence, uncertainty, values, and action. MathGov does not claim perfect access to reality; it requires declared, auditable, challengeable contact with the relevant reality surface.

Root reality doctrine. MathGov defines reality, for decision-making, as the evaluator-independent field of constrained structures that, under transition, produce consequences. Structure is the stable arrangement of relevant differences. Constraint is what limits and enables possible transition. Consequence is what follows when an action, event, policy, or decision enters that constrained structure. In public language: Reality is what exists, what constrains, and what follows.

Normative commitment. Union-Based Ethics states what the framework is committed to protecting and optimizing: rights floors, reduced suffering and domination, catastrophic-risk limits, containment, and flourishing across affected stakeholders. These commitments are explicit rather than hidden inside equations.

Ethical-reality classes. MathGov distinguishes measured or indicator-supported ethical reality, such as injury, mortality, pollution, access, deprivation, or ecological exposure; inferred or uncertainty-bound ethical reality, such as sentience, suffering, agency restriction, psychological harm, or future risk; and explicit normative commitment, such as rights floors, dignity, non-domination, and future-generation standing. Each class requires a different evidence and audit posture.

Formal decision machinery. RippleLogic converts those commitments into gates, scores, weights, audit fields, and decision states. Formal machinery helps make judgment inspectable, but formal consistency alone is not proof that a decision is right.

Evidence and measurement. Impacts, probabilities, sentience estimates, stakeholder maps, and weights require evidence. Where evidence is weak, the framework must disclose uncertainty, run sensitivity checks, escalate, or refuse the stronger claim.

Governance and interpretation. PCCs, public rationales, release manifests, and educational explanations make decisions accountable. They do not create truth by themselves. They make claims easier to inspect, contest, and repair.

Plain-language rule. Do not let one layer pretend to be another. Compliance is not correctness. Plausibility is not evidence. Intelligence is not sentience. Protection is not authority. Scoring is not admissibility. Refusal is not failure; it is a guard against false certainty.

Five-level method. RippleLogic is now best taught inside the MathGov framework as: Reality Grounding -> Rights Floor -> TRC -> CSV -> RLS. Formal shorthand: RG/RSG -> RF/NCRC -> TRC -> CSV -> RLS.

Two-phase teaching frame. The same five-level method is also taught as two phases. **First qualify options:** RG, RF/NCRC, TRC, and CSV determine which options are grounded, rights-safe, tail-safe, and structurally viable. **Second rank survivors:** RLS ranks only the options that survived qualification. This prevents high average benefit from hiding a rights violation, ruin path, ungrounded claim, or structural impossibility.

**Unioning and option generation.** Before accepting a residual trade-off, MathGov seeks cooperative, Pareto-improving, harm-reducing, burden-sharing, staged, reversible, cleaner, or hybrid redesigns. It does not assume that every real conflict can make every affected Union positively better off. Where irreducible conflict remains, protected rights, catastrophic risk, distributional burdens, legitimacy, and justified residual harm remain visible rather than being hidden inside an average score.

Tier-proportional CSV rule. The depth of CSV scales by tier, but CSV does not disappear. Tier 1 uses a brief qualitative screen where the method reaches ranking; Tier 2 requires a proportional, assumption-bounded CSV status; Tier 3 requires full binding CSV evidence. Ordinary RLS ranking is permitted only for pass-equivalent CSV survivors.

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.

Reality Grounding means the decision must say what part of reality it is actually touching before it claims what should be done. A decision sees only a surface, not the whole territory. Evidence grounds the surface. Unknowns mark the edge. Claims must stay inside the ground they stand on.

UCI/HOI placement in plain language. Containment is the structural safety step. UCI and HOI are structural sensors. If the sensors show containment-level damage, that belongs inside Containment. If the remaining options are tied or too close to call after Ripple Score, UCI/HOI may help choose the less brittle, less hollowing option, or the framework may refuse false certainty.

PLSS rule. PLSS is the zoom lens, not the safety waiver. It keeps local decisions practical while preserving rights, tail-risk, containment, and structural-diagnostic duties.

Reality is not whatever the majority believes, whatever an institution certifies, whatever a model outputs, whatever a dashboard computes, whatever a legal authority declares, whatever is easiest to measure, whatever fits the preferred narrative, whatever produces the best score, or whatever the current surface happened to capture. Evidence is not reality. A score is not reality. A PCC is not reality. A model is not reality. Compliance is not correctness. Traceability is not truth. Plausibility is not evidence. Social force is not content-truth. Protection is not authority.

## 2.4 Why MathGov Does Not Claim Perfect Calculation

MathGov does not claim that all values, meanings, risks, harms, probabilities, or social effects can already be measured perfectly. Its purpose is more disciplined: to make assumptions visible, uncertainty explicit, rights floors non-negotiable, catastrophic risks harder to hide, and decisions easier to audit.

A number is not automatically truth. A score is useful only when the evidence behind it is visible, the uncertainty is disclosed, the affected stakeholders can be identified, and reviewers can challenge the

assumptions. When the evidence is not strong enough, RippleLogic should not pretend certainty. It should narrow the claim, escalate for review, run sensitivity analysis, or refuse the stronger decision state.

This is why the system separates admissibility, permission, compliance, evidence, and execution. The goal is not perfect calculation; the goal is better judgment under declared constraints, with less hidden harm and more accountability to the full field of consequences.

Plain-language maturity boundary

The MathGov framework and RippleLogic are not a magic calculator. It does not make ethics automatic, and it does not remove human, legal, democratic, scientific, or professional responsibility. It makes ethical reasoning more structured, inspectable, rights-constrained, uncertainty-aware, and contestable.

Some parts are mature as architecture but still need testing as measurement tools. UCI, biological SGP measurement, calibrated kernels, reference calculators, validators, and empirical pilots remain future maturity work unless separately published and verified.

## 2.5 Reading the Field Without Blurring Its Meaning (Informative)

Scope asks where and at what scale a consequence matters. Dimension asks what kind of condition changes. An EffectToken identifies one distinct consequence. A pathway states how one consequence may cause another. Keep these questions separate. The same injury does not become seven injuries because it matters to a person, a family, a community and a nation.

The dimensions are designed to minimize redundant accounting, not to deny interdependence. Material insecurity can affect health; knowledge can affect agency; belonging can affect meaning. Related is not identical. Seven is the present coverage design, not a proven mathematical limit on human or nonhuman flourishing.

Read the scopes as Self; Family/Household; Community; Organization; Polity; Humanity / Global Coordination, including the CMIU coordination view; and Biosphere. Communities can cross borders, households can include non-cohabiting care, and organizations are not automatically polities. CMIU does not replace people who are not managing anything. The biosphere depends on external cosmic conditions but is not the universe or a light cone.

Raw telemetry is not welfare. A busy server, a signed message or a functioning dashboard does not establish health, experienced meaning, safe execution or successful recovery. Evidence must connect those observations to the actual conditions being evaluated. The matrix preserves the profile; RLS summarizes the qualified residual field. A score is not a percentage of alignment.

A recovery action must work in reality, within rights, risk, capacity and authority constraints. Writing a better number cannot repair a harmed person or ecosystem. Offline fallback, emergency controls and privacy-preserving audit are implementation choices governed by the same method, not extra ethical stages.

## Section 3. Twelve Declared Primitives

A primitive is a foundational starting point that the system does not derive from anything earlier inside the system. A primitive is not a claim of absolute truth. It is a declared foundation, openly stated and open to challenge. The MathGov primitives are organized into three groups: foundational claims about reality (P1 to P4), foundational ethical commitments (P5 to P9), and foundational structural commitments (P10 to P12).

### 3.0 What MathGov Means by Reality

Root reality doctrine. MathGov defines reality, for decision-making, as the evaluator-independent field of constrained structures that, under transition, produce consequences. Structure is the stable arrangement of relevant differences. Constraint is what limits and enables possible transition. Consequence is what follows when an action, event, policy, or decision enters that constrained structure. In public language: Reality is what exists, what constrains, and what follows.

Reality is named here operationally, not metaphysically. Consistent with the Canon's operational UBR boundary, MathGov does not assert a complete theory of being. It asserts something narrower and testable: decisions are embedded in an evaluator-independent structure of constraint and consequence that holds whether or not any agent perceives, admits, models, or has power over 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, or model outputs when they shape welfare, rights, risk, agency, or available choices.

Root-science boundary. MathGov treats constrained structure under transition producing consequence as the operational substrate for governance-relevant reality. This does not claim that reality is only physical matter or that MathGov has final metaphysical access. It means that physical, biological, experiential, social, institutional, informational, artificial-system, and formal/logical structures and constraints become decision-relevant when they shape rights, risk, welfare, agency, available choices, legitimacy, or auditability.

A factor enters the reality surface of a decision when it bears on outcomes, welfare, rights, risk, containment, legitimacy, agency, or the structure of available choices. Constraint has at least two species: causal constraint, such as physical, biological, ecological, institutional, or social effects; and formal/logical constraint, such as coherence, conservation, computability, consistency, and mathematical possibility. MathGov uses both and must not collapse one into the other.

Three distinctions keep this primitive disciplined. First, vehicle, content, and effect must be separated: a belief, norm, institution, model, market signal, or legal status can be a real vehicle with real effects even when its content is false. Consensus can therefore be socially real and causally powerful while still failing to track reality. Agreement does not make a harmful system harmless. Second, surface and territory must be separated: every model captures only a reality surface, while the territory exceeds every surface. Omitted factors are recorded as unknown, not zero. Third, reality, truth, and value must be separated: reality is what holds, truth is a sufficiently faithful map, and value is what matters. P1 fixes the first and grounds, but does not replace, the ethical commitments that follow.

Experience is admitted to the reality surface in its own right, not solely because it may cause outward behavior. Pain, fear, attention, agency, dignity, and welfare-bearing experience are decision-relevant even when access to them is indirect and evidence-governed. SGP protection therefore does not depend on winning a contested claim that experience must be causally efficacious before it can matter.

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. Because catastrophic outcomes are, by the TRC definition, potentially irreversible or unbounded, folding them into an average computes over a tail finite knowledge cannot reliably bound and misrepresents an uncompensable loss as compensable. Finitude and irreversibility motivate the stronger precaution adopted under P6 and CSV; P1 alone does not entail those normative commitments, their exact operators, or their thresholds.

## Group A. Foundational claims about reality

### P1. Reality Constrains Outcomes

Reality is not created by opinion, authority, ideology, wealth, status, institutional consensus, model output, or persuasive language. Those forces may be real vehicles with real effects, but their content can still be false. A policy can sound noble and still cause harm. A profitable system can still degrade the biosphere. A popular leader can still be wrong. A technically impressive AI system can still be misaligned. A social metric can rise while trust, health, ecological resilience, or institutional legitimacy collapses. Decisions have consequences whether or not anyone admits them.

Statement: Reality constrains outcomes, and decision systems must remain accountable to real consequences under declared uncertainty.

## P2. Decisions Create Ripples

Every decision produces effects beyond the immediate actor. Effects may be direct or indirect, short-term or long-term, visible or hidden, local or systemic. This is the foundation of RippleLogic.

Statement: Actions propagate through interconnected systems, producing measurable and partially uncertain ripple effects.

## P3. Stakeholders Exist Within Layered and Intersecting Unions

No entity exists in complete isolation for governance-relevant domains. Persons, institutions, technologies, animals, ecosystems, and future intelligences participate in layered and intersecting patterns of interdependence. The canonical operational stack used in RippleLogic Tier 1 to Tier 3 computation is U1 Self, U2 Household, U3 Community, U4 Organization, U5 Polity, U6 Humanity / Global Coordination, and U7 Biosphere. U6 has two non-identical analytical views: the Humanity view covers all affected human beings and future human populations; the Collective Managing Intelligence Union (CMIU) view covers managing intelligences and institutions participating in shared governance, coordination, and systemic-risk management. Protection and standing do not automatically confer governance authority. Two further layers, U8 Cosmic and U9 Universal / All-Encompassing Infinite Union (AIU), are recognized as non-computational meta-unions and serve only as humility horizons and precautionary amplifiers under the current release line.

Statement: All decision-relevant beings and systems exist within layered, intersecting unions of relationship and consequence. The seven operational scopes are an ordered coverage sequence, not a claim that every real membership relation forms one strict subset chain.

## P4. Welfare Is Multidimensional

Well-being cannot be reduced to money, pleasure, productivity, GDP, or survival alone. The current canonical welfare dimensions are D1 Material, D2 Health, D3 Social, D4 Knowledge, D5 Agency, D6 Meaning, D7 Environment. These are operational categories, not metaphysical absolutes; they are open to refinement under governed update.

Statement: Welfare must be evaluated across multiple dimensions, not collapsed into a single crude metric.

## Group B. Foundational ethical commitments

### P5. Rights Floors Are Non-Compensatory

Some harms cannot be traded away because aggregate benefit appears positive. A system must not be able to say: most people benefit, so this group can be exploited. This is the basis of the Non-Compensatory Rights Constraint (NCRC). Welfare scoring may rank admissible options; it may not override rights floors.

Statement: Certain rights and protections function as floors, not negotiable preferences.

### P6. Catastrophic Risk Cannot Be Averaged

Some risks are too severe, too irreversible, or too civilization-threatening to be treated as ordinary costs inside an expected-value calculation. A small probability of catastrophic failure still matters when the downside is structurally irreversible. This is the basis of the Tail-Risk Constraint (TRC), implemented through Conditional Value-at-Risk (CVaR) over a governed scenario set.

Statement: Catastrophic and irreversible risk must be bounded, not averaged.

### P7. Larger Systems Constrain Smaller Ones

A system can win locally while damaging the larger system it depends on. A company can profit while destroying worker health. A nation can grow while degrading the biosphere. An AI lab can advance capability while raising alignment risk. This is the basis of the CSV gate.

Statement: No system may optimize one layer of reality by destabilizing the larger systems that sustain it.

## **P8. Direct welfare protection is responsive to credible evidence of welfare-bearing experience, subject to precaution, typed protections, and non-downgrade floors.**

A rock, a bacterium, a tree, a forest, a dolphin, an elephant, an orangutan, a human, and a future artificial intelligence do not occupy identical ethical positions. Moral consideration should not be based on species membership, power, familiarity, or economic convenience alone. The Sentience Gradient Protocol v8.8 evaluates protection-relevant moral patienthood through a Moral Patienthood Surface (MPS) centered on valence and welfare evidence, experiential integration, temporal continuity, flexible state-dependent behavior, and substrate or causal support. Full Protection Plateau (FPP) is a separate non-downgrade rights status. Governance Participation Readiness (GPR), Stewardship/Power Readiness (SPR), Intelligence Capacity Profile (ICP), and the informative Reality-Management Capacity Profile/open P100 calibration are assessed separately so that limited agency, language, ethical participation, or reality-management capacity cannot be used to erase credible welfare protection. Every human person has FPP-100 by constitutional anti-domination commitment; humanity collectively is the provisional current RMCP-P100 reference calibration anchor, without implying moral superiority, wise performance, CMIU membership, or authority; non-human individual or class FPP and cross-substrate P100 status follow governed review.

Statement: Moral consideration is responsive to evidence of sentience and welfare-bearing capacity under uncertainty, while every human person retains non-downgradable FPP-100 protection.

## **P9. Domination Is a Distinct Harm**

The Minimal Normative Axiom of Canon Section 3.1 centers not only suffering but domination. Concentrated, unaccountable, or coercive power over others can degrade dignity, agency, and the conditions for flourishing even where measurable welfare appears unchanged. Anti-capture, non-domination, and the preservation of meaningful contestability are first-order concerns of the framework.

Statement: Reducing unnecessary suffering and reducing domination are co-equal normative concerns; neither can be traded for the other.

## **Group C. Foundational structural commitments**

### **P10. Moral Status Is Not Decision Authority**

Four separations are constitutive of the framework and must not be collapsed: (i) moral status is not decision authority; (ii) protection is not governance role; (iii) sentience is not the same as intelligence; (iv) admissibility logic is not welfare optimization. These separations matter because each pairing is regularly conflated in real institutional practice, and that conflation is one of the routes by which rights, audit, and containment are eroded.

Statement: Moral status, decision authority, protection, governance role, sentience, intelligence, admissibility, and optimization are distinct objects and must not be collapsed.

### **P11. Uncertainty Must Be Made Visible**

Most real decisions involve incomplete evidence, model uncertainty, value disagreement, measurement error, unknown side effects, and strategic behavior. The danger is not uncertainty itself; the danger is hidden uncertainty pretending to be certainty. RippleLogic requires uncertainty to be represented openly through declared methods, sensitivity analysis, and audit-flag semantics.

Statement: Uncertainty must be represented openly rather than buried inside authority, rhetoric, or false precision.

### **P12. Power Must Be Auditable**

Human authority alone is not stable enough at civilization scale because humans and institutions can be corrupted, coerced, captured, bribed, threatened, gamed, or distorted by incentives. The remedy is not to remove humans from governance; the remedy is to constrain authority by transparency, evidence, declared procedure, and contestability. The Provenance and Compliance Certificate (PCC), the Five-Sentence Public

Rationale (5SPR), and the Notice-Choose-Act-Reflect (NCAR) loop exist to make legitimate authority inspectable.

Statement: Legitimate decision power must be inspectable, challengeable, evidence-linked, rights-constrained, and accountable over time.

## Section 3A. Primitive Dependency Audit

The twelve MathGov primitives are declared starting points for this framework. They are not claims of absolute metaphysical irreducibility. Each primitive depends on background concepts, evidence traditions, measurement practices, and normative choices that must remain visible.

This matters because a decision framework becomes more trustworthy when it declares where it begins, what it imports, what it assumes, what it measures, and what could force revision.

The following audit posture applies to all twelve primitives.

1. A primitive is a declared starting point inside MathGov, not an assertion that no deeper analysis is possible.
2. A primitive may be imported from parent disciplines such as systems science, ethics, law, risk science, cognitive science, ecology, political theory, or decision theory.
3. A primitive remains legitimate only if its dependencies are disclosed and its operational use remains testable, contestable, and revisable under governed update.
4. A primitive should not be protected by reinterpretation after reality contradicts its projected consequences.
5. Where a primitive cannot generate a unique operational consequence for a run, the correct response is not false certainty. The correct response is refusal, escalation, or bounded provisional action under declared uncertainty.

Primitive dependency checklist:

For each primitive P1 through P12, reviewers should ask:

- What does this primitive presuppose?
- What parent disciplines does it import from?
- What operational role does it play in RippleLogic?
- What would count as evidence that this primitive is incomplete, misframed, or misapplied?
- Does this primitive affect admissibility, scoring, interpretation, or public communication?
- Does the current run depend on this primitive in a way that should be disclosed in the PCC?

Plain-language rule. MathGov does not become stronger by pretending its primitives float above all prior thought. It becomes stronger by declaring its starting points, exposing their dependencies, and letting reality test their consequences.

Summary statement. The MathGov primitives are declared foundations for an ethical decision architecture. They are not final metaphysical proof. They are disciplined commitments that must remain visible, testable, and open to governed refinement.

## Section 3B. The Reference Stack: Against What Is the Decision Being Judged?

A decision system is not trustworthy simply because it has rules, dashboards, audits, or guardrails. Those tools matter only when they are connected to a valid reference.

RippleLogic asks a simple but demanding question: Against what is this decision being evaluated?

For a bridge, the reference may include load, stress, gravity, material strength, and failure thresholds. For governance, the reference is more layered. It includes bodies, needs, rights, ecosystems, institutional legitimacy, causal evidence, uncertainty, and audit records.

RippleLogic does not pretend that every governance question has one simple physical reference. Instead, it uses a reference stack: physical constraints; biological needs; rights floors; catastrophic-risk corridors; containment conditions; stakeholder impacts; governance authority; evidence status; and auditability.

This is why The MathGov framework and RippleLogic are not merely a scoring system. It is a reference-structured decision architecture. If the reference is missing, circular, unfalsifiable, or too weak for the claim being made, the honest answer is not forced selection. The honest answer is refusal, escalation, more evidence, or a narrower claim.

Plain-language rule. Do not govern by assumption. Declare the reference, show the evidence, disclose the uncertainty, and refuse false certainty.

## Section 3C. Category Grounding: Are the Words Still Attached to Reality?

Reality Grounding asks what part of reality the decision has actually touched. Category Grounding asks whether the categories used to describe that reality are bounded correctly enough for the claim being made.

This matters because a wrong category can make a good-looking metric dangerous. If the category is wrong, the metric may measure the wrong thing. If the metric measures the wrong thing, optimization can improve the wrong target. If governance protects the optimized metric, the system may become increasingly efficient at maintaining drift.

Simple example. If a school optimizes "learning" but defines learning only as test-score improvement, it may raise the metric while weakening curiosity, confidence, teacher trust, long-term understanding, and student agency. Category Grounding asks whether "learning" was defined correctly before the score is optimized.

Category Grounding is not a sixth step. It belongs inside Reality Grounding when material. The five-level method remains: Reality Grounding -> Rights Floor -> TRC -> CSV -> RLS.

A category is material when it affects who counts as a stakeholder, what counts as harm, what counts as sentience or agency, what counts as safety or stability, what counts as catastrophic risk, what counts as containment damage, what counts as authority, or what metric will be optimized.

Reviewers should ask: What does this term mean? What boundary defines it? What belongs inside it? What does not belong inside it? What discriminator separates it from nearby terms? Where does it sit in the dependency order? What reference structure tests it? What evidence surface supports it? What would force revision or refusal? What happens if this category is wrong?

Plain-language rule. Before MathGov scores a thing, it checks whether the thing has been named, bounded, and grounded well enough for the decision. A score is not a cure for a bad category.

Semantic debt. Semantic debt is the cost of continuing to use categories whose reality correspondence has become weak, stale, inherited, contested, or unreviewed. MathGov does not become stronger by hiding semantic debt. It becomes stronger by naming it, narrowing claims, and revising categories when reality requires it.

## Section 3D. Structural Viability: Can the Plan Actually Stand?

After a decision is grounded, rights-safe, tail-risk bounded, and containment-safe, one more question may matter before scoring or execution: can this plan actually stand? Structural Viability Check asks whether the option has the resources, dependencies, reversibility or endpoint clarity, operational capacity, and internal consistency required for current execution.

This is not a sixth public step. It belongs inside Containment when execution feasibility is material. The public method stays: Ground -> Protect -> Bound -> Contain -> Score.

Plain rule: a good intention is not enough if the plan cannot structurally work. But structural viability is not justice either. A plan can be stable and still violate rights. Rights floors remain first and non-negotiable.

A structurally non-viable option may still be kept as a redesign target. It should not be presented as currently selectable, executable, deployment-ready, or deterministically selected.

## Section 3E. No Silent Remainder: Unknown Is Not Zero

MathGov allows uncertainty. It does not allow silent blanks. A required field may be estimated, assumed, declared unknown, not applicable, or blank by design. It may not simply disappear. Silent gaps are how false certainty enters a decision.

Plain rule: if the decision depends on it, name it. If you do not know it, say you do not know it. If it does not apply, say why. Do not turn unknown into zero.

## Section 4. Union as a Reality Principle

Union does not mean everything is the same. In MathGov, it means that decision-relevant beings and systems exist through relationships, dependencies, and consequences.

This is the core of what Union-Based Reality (UBR in the canonical specification) describes. A human being is distinct, but never isolated. The body is built from organs, cells, water, minerals, breath, food, ancestry, language, culture, and environment. A family is built from individuals. A society is built from families, institutions, laws, ecosystems, and shared meanings. A biosphere is built from interdependent life. Every intelligence or intelligence-like system we have actually observed, biological, social, or artificial, operates through relationships among memory, perception, feedback, environment, and context. Within governance-relevant empirical domains, observed biological, social, ecological, and artificial systems depend on relations, inputs, constraints, and environments. RippleLogic therefore does not assume complete operational isolation; it treats any isolation claim as bounded and evidence-dependent rather than as a universal metaphysical impossibility theorem.

That is what union means in this framework. Not a doctrine, not a slogan, not a feeling. It is the operational condition that decisions happen inside structured relationships, not inside isolated abstractions.

Union does not erase difference. Water is not fire. A human is not a dolphin. A tree is not a government. A child is not an artificial intelligence. Distinctions matter. Boundaries matter. Rights matter. The framework's commitment to non-compensatory rights floors only makes sense because individual beings are real and their integrity is real.

But distinction is not absolute separation. Every being and system exists within larger fields of dependency and consequence. Actions ripple. Care propagates. Harm propagates. Trust propagates. Collapse propagates.

These do not propagate at the same rate. Reality is not symmetric. A drop of poison ruins clean water; a drop of clean water does not redeem poisoned water. Construction takes years and destruction takes minutes. Trust is built slowly and lost quickly. Health degrades faster than it heals. A lie travels around the world before the truth has its boots on. Union is not balance in the sense of equal harm and equal benefit. Reality is asymmetric, and any honest system has to say so.

This rules out three common misreadings of union, all of which MathGov rejects.

The first is cosmic merger: everything is one, so distinctions and harms do not really matter. This dissolves rights, accountability, and the moral weight of suffering, and it usually arrives wrapped in beautiful language. The framework rejects it.

The second is required opposition: every good must be matched by an equal bad to keep things in balance. This is a misuse of yin and yang, of karma, of dialectic, and of the rhetoric of "both sides." A person does not need to do harm to do good. A system does not need to produce equal damage and benefit to be healthy. Clean water sustains life better than poisoned water. Truthful information serves people better than deception. Health is better than preventable sickness. Cooperation is more stable than domination over the long run. Better balance does not mean equal parts harm and benefit. Better balance means right relationship.

The third is collectivist dissolution: the individual is only a fragment of the whole, so individual rights bend to collective will. This is the totalitarian misuse of unity, and it is precisely the failure mode the rights gate (NCRC) is built to refuse.

Union also does not flatten attention. A breakfast choice acts strongly on the self, weakly on the household, and almost not at all on the polity. A national infrastructure decision acts the other way around. The discipline is to match moral and analytical attention to actual coupling, dependency, and consequence, not to perform equal concern at every scale.

So the MathGov claim is not "everything is one, therefore nothing matters." The claim is the opposite. Things have real identities, and those identities exist inside networks of cause, dependence, and effect. Decisions have to be evaluated by their ripples across the systems they actually touch.

Different vocabularies converge on the same observation. Things are connected. Connection has structure. Structure has consequences. Consequences are not symmetric. Decisions matter.

Union is not dogma. In the operational sense used here, union is the condition that every decision unfolds inside a field of relationship, dependency, consequence, and accountability.

## Section 5. Eleven Declared Assumptions

An assumption is a provisional claim used to make the system runnable. Assumptions are weaker than primitives. They must remain visible, challengeable, and revisable as evidence and practice accumulate.

### A1. The seven welfare dimensions are a strong starting set

Material, Health, Social, Knowledge, Agency, Meaning, and Environment appear broad enough to capture most major welfare effects relevant to current governance. They may be refined, subdivided, or expanded under governed update. They are not claimed to be metaphysically final.

### A2. The seven operational unions are sufficient for most current decisions

U1 Self through U7 Biosphere are sufficient for most contemporary governance, policy, business, AI, and ethical decisions. The Cosmic and Universal layers should remain humility horizons, not over-quantified before evidence supports parameterization.

### A3. Rights floors can be specified well enough to guide decisions

Societies can identify core protected domains: life, bodily integrity, basic needs, dignity, liberty, due process, information integrity, and ecological integrity. Disagreement at the edges does not erase the reality of obvious violations; the existence of hard cases does not eliminate clear ones.

### A4. Sentience can be evaluated on a gradient under uncertainty

Protection evidence is not all-or-nothing, but it must not be collapsed into one prestige score. SGP v8.8 reports a theory-plural MPS band and interval, process confidence, protection posture, and FPP status. Valence and welfare evidence are primary; higher cognition, agency, communication, intelligence, and ethical participation may support interpretation but are not prerequisites for suffering. Human FPP is absolute, while non-human protection is precautionary, evidence-governed, and open to full plateau recognition.

## **A5. Weights can be made legitimate through Hybrid Democratic Weighting**

Welfare weights become more legitimate when they combine constitutional floors, stakeholder input, expert evidence, democratic review, sensitivity testing, and public audit. Hybrid Democratic Weighting (HDW) is the canonical mechanism for this combination. PLSS, where declared, governs prominence-adjusted residual welfare ranking under the same anti-capture logic.

## **A6. Quantification can improve ethics if it remains humble**

Not everything valuable can be perfectly measured. Structured estimation is often better than unstructured power. The danger is false precision; scores must be treated as decision aids, not divine commands. RippleLogic addresses this through layered cascade, audit flags, and the prohibition of Tier-4 claims absent ProofPack.

## **A7. Open systems are less capture-prone than closed systems**

Open standards, public documentation, audit trails, and reproducible decision records are intended to reduce capture risk. Whether they reduce manipulation or concealment in practice requires empirical comparison; openness alone does not establish integrity.

## **A8. Better decision architecture can reduce capture and authoritarianism**

MathGov does not claim to defeat authoritarianism, corruption, or destabilizing power concentration. It aims to make these failure modes substantially harder to sustain by forcing decisions into visible structures: who is affected, what rights are at stake, what harms are hidden, who benefits, who bears the risk, what evidence supports the decision, and whether the reasoning can be reconstructed later.

## **A9. AI and digital systems will increasingly participate in governance**

Whether or not current AI is sentient, AI systems already influence decisions: ranking information, automating workflows, recommending actions, detecting patterns, shaping perception, and mediating institutional behavior. Governance architecture must account for AI as both a tool used by decision-makers and a possible future stakeholder if SGP-grade evidence of moral patienthood emerges.

## **A10. The system itself can be gamed and must model agents**

Any governance method can be captured, bureaucratized, moralized, weaponized, ignored, ritualized, or turned into ceremonial rationality. The framework must therefore model agents and incentives: who benefits, who resists, who can manipulate the process, who can route around it, and who bears the risk when the system fails.

## **A11. Corrigibility must hold, with a floor on what is revisable**

MathGov must remain corrigible, evidence-responsive, and open to structural improvement, otherwise it ossifies into dogma. At the same time, corrigibility itself requires a floor: the rights gate, the catastrophic-risk gate, the CSV gate, and the auditability commitment cannot be revised away by ordinary stakeholder process. Revision of these floors is constitutional in character and requires governed extension protocol, not ordinary update.

# **Section 6. What MathGov Does Not Assume**

Equally important to what MathGov declares is what it refuses to claim. The framework explicitly does not assume:

- That it possesses final ontology or final moral truth.
- That all values reduce perfectly to numbers.

- That humans should surrender governance to AI systems.
- That current AI systems are conscious or moral patients.
- That all beings have equal moral status.
- That powerful beings matter more than vulnerable beings.
- That a majority can vote away fundamental rights.
- That data is neutral or self-interpreting.
- That models cannot be captured.
- That formalization eliminates moral conflict.
- That one culture, nation, species, or institution should define flourishing for all others.
- That a Tier-4 conformance claim is currently legitimate (it is prohibited until ProofPack is publicly replayable and independently verified).

This boundary statement exists because the framework must avoid becoming exactly what it is designed to resist: a rigid authority system pretending to be universal while concealing its own assumptions.

## Section 7. The Operational Cascade

RippleLogic operates as a five-level admissibility-and-ranking method. Each level answers one distinct question. Levels do not collapse into a single blended score; later levels do not override earlier levels; earlier success does not excuse later failure.

Level 1, Reality Grounding / Reality-Surface Grounding. What reality surface, evidence trace, material unknowns, category boundaries, transition boundary, consequence pathways, and claim boundary support the run? If the claim exceeds its reality surface, the run must narrow, escalate, gather evidence, mark the result exploratory, or refuse the stronger claim.

Level 2, Rights Floor / NCRC. Does this option cross a non-compensatory rights minimum? An option that does is inadmissible. Welfare ranking cannot rescue a rights-floor violation.

Level 3, Tail-Risk Bound / TRC. Among rights-admissible options, is the catastrophic downside acceptably bounded? An option whose CVaR over the governed scenario set exceeds the declared corridor is inadmissible, regardless of expected welfare.

Level 4, CSV / Containment and Structural Viability. Among admissible options, does this option preserve sufficient integrity of the larger system that contains the local gain? An option that fails Containment is not selectable, even if it ranks highly on residual welfare.

Level 5, Ripple Score / RLS. Among selectable options, which produces the best residual welfare pattern under the declared weighting regime? RLS ranks; it does not rescue failures in Reality Grounding, NCRC, TRC, or CSV.

Special-case UCI / HOI tie-break. UCI and HOI are structural diagnostics used inside Containment when material. Only after options are already in the selectable set may residual UCI/HOI help document a tied, close, uncertainty-overlapped, or non-decisive RLS result. They do not reopen earlier gates and they are not a substitute welfare score.

The non-overlap rule is canonical: Reality Grounding = claim authority; NCRC = rights feasibility; TRC = catastrophe feasibility; Containment = cross-scale structural feasibility; RLS = residual welfare ranking within the selectable set; residual UCI/HOI = non-decisive structural resolution or monitoring within the selectable set. No implementation, commentary, or downstream artifact may collapse these functions into a single blended score.

## Section 8. The Five Decision States

RippleLogic recognizes five distinct decision states. They must not be collapsed, and a later state must not be used to infer an earlier one.

**Rights-admissible.** The option passes the Rights Floor / formal NCRC. Reality Grounding still controls the strength and scope of every claim made about that result.

**Admissible.** The option is supported within its declared Reality-Grounding claim boundary (RG\_SUPPORTED or RG\_NARROWED), passes RF/NCRC, and has a qualifying Tail-Risk status: TRC\_PASS or a documented TRC\_NOT\_TRIGGERED. Admissibility is not yet selectability.

**Selectable.** The option is admissible and has a qualifying CSV status: CSV\_PASS, CSV\_PASS\_WITH\_CONTROLS, or a properly documented CSV\_NOT\_MATERIAL for the declared claim boundary.

**Selected.** The option is chosen from the selectable set. Record whether this is framework selection under the Canon's selection conditions or separate authority selection after a non-decisive result. A tie-break preference alone does not remove framework refusal or grant execution authority.

**Refusal.** The framework cannot uniquely determine the applicable admissibility or selection state from the declared inputs, evidence, projection rules, or layer commitments.

No downstream score may rescue failure at an earlier state. A welfare gain cannot make a rights-violating option admissible. A high RLS cannot make a containment-failing option selectable. Refusal must not be silently converted into gate failure, gate success, or discretionary selection. The decision-state vocabulary exists precisely to prevent these substitutions.

### 8.1 What refusal means in plain language

Refusal means the framework cannot honestly make the claim being asked of it. It is not the same as rejecting an option on ethical grounds. It is not the same as selecting the safest-looking option by instinct. It is a statement that the available evidence, declared projection, or category structure is insufficient for the requested verdict.

For example, if a decision lacks enough evidence to determine whether a protected subgroup is harmed, the correct result is not to assume safety. If a scenario model cannot determine whether catastrophic risk is inside the corridor, the correct result is not to average the uncertainty away. If a sentence claim relies only on fluent language or self-report, the correct result is not to treat that as proof of moral patienthood.

Refusal can lead to more evidence gathering, redesign, escalation, a narrower claim, or bounded provisional action where delay itself would create harm. What it may not do is pretend that uncertainty has been solved.

## Section 9. Critical Separations

Four separations are constitutive of the framework and merit explicit restatement, because each is regularly conflated in institutional practice and each conflation opens a route to rights, audit, or containment failure.

**Moral status is not decision authority.** A being can warrant moral consideration without thereby acquiring governance role. Authority is grounded in responsibility, accountability, and procedural legitimacy, not in moral status alone.

**Protection is not governance role.** A protected stakeholder is not, by virtue of protection, a decision-maker over others. SGP confers grounds for protection; it does not confer authority.

Sentience is not intelligence. Capability and welfare-bearing experience are distinct. A highly capable system may not be a moral patient; a moral patient may have limited intelligence or agency. SGP v8.8 uses MPS/FPP for protection and keeps the optional Intelligence Capacity Profile, GPR, and SPR separate.

Admissibility is not optimization. RF/NCRC, TRC, and CSV are feasibility gates. RLS is welfare ranking. They have different inputs, different operators, and different failure consequences. A welfare maximum that is not admissible is not selectable; an admissible option with mediocre RLS is still selectable.

## Section 10. SGP v8.8 Integration: Protection, Participation, Reality Management, and Power

SGP v8.8 is the canonical companion specification for moral patienthood and rights-of-protection under uncertainty. It replaces the former minimum of Awareness, Agency, and Union Participation with an explicit separation of constructs:

- **MPS** asks how strongly evidence supports welfare-bearing experience.
- **FPP** establishes full non-downgradable protection.
- **GPR** asks how an entity can participate in a named governance process.
- **SPR** asks whether an entity is eligible for a named consequential role under constraint.
- **ICP** records intelligence and understanding capacities without treating them as sentience or authority.

### 10.1 Human Full Protection Plateau

Every human person has FPP by equal-dignity and anti-domination commitment. This is not an empirical claim that every human has maximal measured consciousness or intelligence. Infants, sleeping or anaesthetized persons, people with dementia, severe disability, communication differences, illness, or temporary incapacity remain fully protected. Their RippleLogic welfare-inclusion coefficient is fixed at  $h_{\{k,welfare\}}=1$ ; this is a constitutional inclusion rule, not an MPS evidence interval.

Non-human biological or artificial individuals and classes may also receive FPP through governed SGP review. Language, human likeness, reciprocal morality, or governance competence are not prerequisites for full protection.

### 10.2 MPS interval binding: residual Welfare stream only

For a valid non-FPP SGP record, RippleLogic uses the MPS evidence record to define the welfare-inclusion hypotheses tested in the residual Welfare stream. The MPS interval is not inserted as a cardinal multiplier. If the preferred option changes across the governed hypothesis set, the run must disclose

**MPS\_HYPOTHESIS\_SENSITIVE**, seek more evidence, adopt a robust precautionary policy, narrow the claim, preserve multiple options, or refuse a definitive ranking.

The Base stream used by RF/NCRC, TRC, and CSV fixes the SGP factor at 1. SGP cannot lower a rights floor, soften a catastrophe bound, weaken containment, or turn uncertainty into zero.

### 10.3 Participation and authority remain separate

GPR may support direct, assisted, or represented participation. A dog, infant, disabled person, octopus, or possible digital mind may deserve protection even when direct deliberation is impossible. Representation is not exclusion.

SPR is role-specific. It requires Reality Grounding, domain competence, rights reliability, tail-risk control, containment and structural viability, corrigibility, auditability, non-domination controls, lawful mandate, and revocability. Neither MPS, FPP, GPR, nor intelligence grants power by itself.

## 10.4 CMIU belongs to governance, not moral-status scoring

The Collective Managing Intelligence Union is the Canon's U6 governance-coordination view: intelligences, representatives, and institutions with a declared participation channel or lawful role in shared governance, coordination, and systemic-risk management. It is not a sentience tier or protection plateau. FPP does not create CMIU membership; GPR informs participation mode; SPR and lawful mandate govern consequential roles.

## 10.5 Minimum integration record

An SGP-aware run records the target boundary, SGP version and hash, MPS band or MPS-NE status, evidence-support interval where applicable, confidence, protection posture, FPP, evidence trace and dependence clusters, reviewer status, alternative explanations, welfare-inclusion hypotheses tested, and revision trigger. Missing, poorly observable, or inadmissible data are not converted to MPS-0 or zero.

## 10.6 Human dual anchor and open RMCP-P100 reference plateau

SGP v8.8 uses 100 in two explicit, non-collapsed ways. Every human person holds FPP-100, the full and non-downgradable protection plateau. Humanity as a collective is the provisional current RMCP-P100 reference calibration anchor because it combines cumulative symbolic culture, formal mathematics and science, intergenerational knowledge, technological intervention, institutions, planetary monitoring, cross-scale causal reach, and Recursive Union Awareness.

P100 is open rather than species-exclusive. Another biological, digital, hybrid, collective, or extraterrestrial intelligence may reach, exceed, or already exist at that plateau when admissible evidence supports functional equivalence or greater capacity. RMCP/P100 is not evidence of sentience, moral worth, wise stewardship, CMIU membership, lawful mandate, or authority. A highly capable system may have low or unknown MPS; a strongly protected being may have limited RMCP. Greater capability increases responsibility rather than permission to dominate or convert protected beings into resources.

### Section 11. The Operating Sequence

A RippleLogic decision proceeds through eleven steps. Tier requirements (Section 15 of Canon) determine which steps are mandatory at which level of rigor.

1. Define and Reality-Ground the decision. State the decision, decision-maker, tier, configuration, baseline, candidate options and no-action condition. Record the reality surface, evidence trace, material unknowns, transition boundary, consequence pathways and supported claim boundary.
2. Map stakeholders across affected unions. Identify the relevant scope coverage sequence from U1 Self through U7 Biosphere; note where U8 Cosmic or U9 Universal / All-Encompassing Infinite Union (AIU) humility considerations apply.
3. Identify sentient and welfare-bearing stakeholders under SGP. Apply governed SGP procedure for class-level and entity-level moral status and rights of protection.
4. Apply the Rights Floor (formal NCRC). Check the floor, categorical-prohibition and severe-hazard channels, including the worst-off protected subgroup and full unallocated rights effects. A failed or materially unresolved rights check cannot pass through an average welfare gain.
5. Apply TRC. For each rights-admissible option, evaluate CVaR over the governed scenario set against the declared corridor. Eliminate any option that exceeds the corridor; record Tail Emergency Mode if invoked.
6. Apply CSV: Containment and Structural Viability. Review the qualified option's controls, dependencies, resources, authority, monitoring, reversibility, execution feasibility, post-state viability and future safe option space. Containment Mode A is a subtest where required, not the whole CSV gate.
7. Score residual welfare via RLS. Rank selectable options across the active 49-cell welfare matrix using HDW (or PLSS, where declared) weights and declared uncertainty method.

8. Make weights, evidence, and uncertainty visible. Report the ex-ante locked union and dimension weights, evidence quality, stakeholder input, expert adjustments, uncertainty treatment and sensitivity tests in the PCC; this reporting step does not authorize post-result parameter choice.
9. Model agents and incentives. Identify who might resist, game, distort, or route around the decision; record relevant audit flags and mitigation plans.
10. Produce an auditable decision record. Emit the Provenance and Compliance Certificate (PCC) with the Five-Sentence Public Rationale (5SPR) and any required containment, SGP, and stewardship disclosures.
11. Run NCAR. Notice, Choose, Act, Reflect. After implementation, evaluate outcomes, surface missed harms, revisit weights, and trigger governed update where evidence requires it.

## Section 11A. A Simple Reading of the System

A simple act can still be read through the full architecture without overcomplication. Suppose a parent has money to buy a toy for their own child and notices another child nearby who has less. The parent could buy for their own child, buy for the other child, buy something shareable, save the money, or do nothing.

MathGov does not command one sentimental answer. It asks for disciplined attention. Are any rights at stake? Is anyone being coerced, neglected, humiliated, or endangered? Is there any catastrophic risk? Usually not. Is there a containment issue, such as undermining the parent's household obligations or creating unhealthy dependency? Usually minor, but still worth noticing. Then the decision moves into residual welfare: self, household, nearby community, social trust, meaning, and material cost.

In a low-stakes moment, the full calculation can be mental and lightweight. In a high-stakes policy decision, the same logic becomes formal, documented, and auditable. This is the point of tiers: the structure is stable, but rigor scales with consequence.

The deeper lesson is that RippleLogic is not only for governments or AI systems. It trains attention. It asks: who is touched, what is protected, what could go wrong, what improves, what is uncertain, and what should be learned after action?

## Section 11B. Qualification Continuity in Plain Language

A decision is not finished merely because an option was scored and selected. Before a consequence-bearing action becomes real, MathGov asks whether the *same option*, *same action*, *same configuration*, *same evidence*, *same controls*, and *same authority conditions* still hold.

That is qualification continuity. A past pass cannot authorize a changed present. If the model, tool, permission, environment, affected people, evidence, controls, or authority changes materially, inherited qualification pauses and the affected parts of the run reopen.

The order remains clear. Reality Grounding, rights protection, catastrophic-risk bounds, containment, and selection come first. Qualification continuity does not add a gate and does not grant authority. It binds the selected option to an exact action and checks that the action remains inside the qualified envelope.

After action, MathGov compares observed outcomes with predeclared markers. New harm, control failure, material deviation, or a new unknown triggers requalification. A good observed result supports only the bounded claim justified by its population, duration, evidence quality, and configuration.

Plain-language rule: do not ask only whether this option passed once. Ask whether this exact action is still qualified now—and what evidence will make us pause, repair, or change course after it becomes real.

Simple example. An institution selects a bounded AI-assistant pilot. Before execution it freezes the model and tool versions, permissions, affected user group, monitoring, shutdown, rollback, controls, qualification snapshot, and separate accountable authority. A permission expansion or new affected group forces requalification. After the pilot, observed rights, workload, error, recourse, and containment signals determine whether the run continues, changes, or stops.

## Section 12. Posture: Structured Humility Under Reality

A decision is better when it protects rights, bounds catastrophic risk, preserves the integrity of containing systems, improves welfare across affected unions, makes uncertainty visible, preserves auditability, and remains open to correction.

More compactly: choose the action that creates the best defensible ripple across affected stakeholders without violating rights and without destabilizing the systems that make life possible.

The deeper posture of MathGov is not final truth by declaration. It is structured humility under reality. The framework declares its primitives, exposes its assumptions, protects its rights floors, models uncertainty, documents reasoning, and remains open to correction as better evidence and better models emerge.

The goal is not perfect morality by calculation. The goal is to make ethical and governance decisions harder to corrupt, harder to hide, harder to capture, and easier to test against reality, so that intelligences (human, institutional, and eventually artificial) can align with reality, with rights, with long-term survival, and with the flourishing of interconnected life itself.

## Section 13. Carried-Forward v12.1 Rights and Ruin Hardening Within Current v13.0

Current completion note. Core Release 2026.09 v13.0 preserves the five-stage cascade while carrying forward four operational hardenings: categorical and severe rights exposure cannot be erased by short duration or low confidence; all-TRC-fail cases enter governed Tail Emergency Mode rather than ordinary ranking; RLS is normalized over active effective weight mass; and Method B uncertainty is computed from pre-cancellation contribution mass. The v12.0 Physical Admissibility, Reality Grounding, Term Discrimination, and Semantic Stability boundaries remain in force, together with the v13.0 rights non-dilution and robust-selection controls.

The v10.1–v10.8 historical methodology hardening does not change the soul of MathGov or the core RippleLogic cascade. It makes the scientific posture stricter. A high-stakes claim should not merely sound coherent, fit the past, or be compatible with observed outcomes. It should declare what would count against it before the evidence is used to protect it.

The plain-language rule is simple: say what you are assuming, say what follows if the assumption is right, say what we should observe, say how much error is tolerated, and say when the claim must be refused, downgraded, or abandoned.

### 13.1 The seven-step trace

Primitive commitment → Operator → Regime → Projection → Observable consequence → Tolerance → Verdict.

### 13.2 What this prevents

This prevents post-hoc rescue, category collapse, symbolic overconfidence, vague claims of validation, and the habit of calling a useful model proven. It also prevents the opposite error: refusing to act where a bounded, rights-preserving, reversible, and well-audited provisional action is justified.

### 13.3 What this does not claim

The carried-forward v10.8 method posture does not claim final truth, perfect objectivity, complete prediction, or immunity from formal limits. It says something more practical: every serious claim must carry its own audit trail, evidence boundary, and reality-facing test.

## 13.4 Binary Gates and Uncertain Evidence

MathGov and RippleLogic distinguish two ideas that are often confused. A decision gate can be binary, while the evidence used to evaluate that gate can still be uncertain.

For example, a rights floor may pass or fail. A tail-risk corridor may be exceeded or not exceeded. A containment relation may hold or fail. But the evidence used to evaluate these option-rejecting gates can involve measurement error, incomplete models, probability estimates, contested data, or human judgment.

The carried-forward v10.8 rule is simple: do not pretend uncertainty is gone merely because the decision surface needs a yes/no answer. Binary gates make governance accountable. Evidence humility keeps the binary answer honest.

When evidence is not strong enough to support the requested claim, RippleLogic should not force a confident decision. It should narrow the claim, request more evidence, escalate to review, or refuse the stronger claim. Refusal is not failure; it is the system refusing to turn uncertainty into false certainty.

This is why high-stakes decisions must identify the transition being proposed, the affected scopes, the primitive assumptions, the projection path, the observable consequences, and the conditions that would trigger rollback or refusal before consequence-bearing execution.

## Section 14. What Is Mature Now, and What Still Requires Validation?

Current specification maturity note. Canon v13.0 and SGP v8.8 develop the previously reviewed specification lineage. The current manifest identifies the exact companion files and Reports/Verification\_and\_Readiness.md identifies executed checks and remaining limits. Biological SGP measurement remains guidance-stage rather than measurement-complete; empirical, deployment and ProofPack claims require separate evidence.

RippleLogic is substantially specified as a proposed Tier 1-3 research and teaching architecture. Measurement validity, implementation conformance, empirical reliability, comparative decision value and production assurance are not established by specification completeness.

### 14.1 Specified now

The cascade architecture is specified: Reality Grounding, RF/NCRC, TRC, CSV, then residual RLS ranking. UCI/HOI are structural diagnostics within their declared CSV, monitoring or special-case tie-break roles, not an additional cascade level.

The claim-boundary language is explicit: GitHub-ready does not mean ProofPack-ready, Tier 4-ready, or empirically validated.

The audit vocabulary is strong enough for worked-run disclosure and reviewer challenge.

### 14.2 Still requiring validation

UCI is load-bearing, but cross-domain validated UCI instruments are not yet published.

Kernel propagation is design-ready, but domain-calibrated and backtested kernels remain future work.

SGP biological evaluation is protection-useful, but individualized biological measurement requires species-appropriate annexes and expert validation.

HDW can expose weights, but institutional capture must be actively audited rather than assumed away.

PLSS keeps local decisions practical, but escalation triggers must be visible so local convenience does not hide non-local harm.

### 14.3 Plain-language rule

The framework may be structurally ready before every instrument is fully validated. Good governance requires both: use the architecture now where appropriate, but disclose the maturity of the evidence, measurements, kernels, institutions, and escalation mechanisms behind each claim.

### 14.4 Cross-Artifact Concept Map

RippleLogic Canon v13.0 is the governing artifact. It defines the cascade, equations, admissibility gates, PCC requirements, claim boundaries, conformance vocabulary, reality-grounding records, schema ownership boundaries, and machine-readable registry expectations.

SGP v8.8 is the companion moral-status and protection protocol. It provides typed protection and role evidence objects, including an MPS evidence band/interval or MPS-NE status, FPP, GPR, SPR, optional ICP, and informative RMCP/P100. It does not provide a cardinal moral-worth scalar, grant governance authority, or override RippleLogic admissibility gates.

ripple.md v5.8 is the portable alignment contract and assurance wrapper. It helps implementations record evidence, decision notes, update legitimacy, reality-grounding fields, refusal conditions, and deployment assurance. It does not replace the Canon.

RippleLogic Agent System v13.0 governs agent deployment, security posture, mode control, operator authority, audit logging, reality-reference closure, refusal, rollback, and runtime conformance. It does not authorize an agent to bypass the Canon.

Ripple Aligners Sheet v5.9 is a worked-run exemplar and training companion. It helps users inspect and replay a sample run, including FrameworkVerdict and RealityGroundingStatus surfaces. It is not a ProofPack, semantic validator, or reference calculator.

Foundations Primer v4.7 is explanatory and carries no normative force. It helps readers understand MathGov and RippleLogic in plain language.

Release-integrity files document the release package. They verify file integrity and packaging claims, but they do not provide empirical validation or ProofPack readiness.

The implementation release requires a limited run-record schema, semantic contradiction validator, fixtures, and CI checks for R0/R1 procedural conformance. Their actual availability and successful execution must be established from the supplied implementation inventory; the Core 15 document bundle alone does not establish them. The complete PCC, Decision Note, gate, SGP, Appendix R, reference-calculator, and calibration ecosystem required for ProofPack or Tier 4 remains a separate requirement. Machine-readable tooling improves checkability and replay readiness but does not by itself establish evidence truth, empirical validity, legal authority, deployment safety, or moral truth.

### 14.5 What Changes in Practice

A RippleLogic-informed decision process changes institutional practice in five concrete ways.

First, rights floors are checked before welfare optimization. A high-scoring option cannot rescue a rights violation.

Second, catastrophic tail risk is checked before ordinary ranking. Expected benefits cannot erase unacceptable catastrophic exposure.

Third, local gains must be tested against containing systems. A decision that helps a local actor while degrading a household, institution, polity, humanity, or biosphere-level support system requires escalation or refusal.

Fourth, uncertainty is not hidden. Weak evidence, missing indicators, unresolved maturity limits, and underdetermined projections must be disclosed.

Fifth, the decision record becomes auditable. The system records what was considered, what was excluded, which evidence was used, which assumptions were made, which gates passed or failed, and why selection or refusal occurred.

## Section 15. Practical Validation Path and Civilizational-OS Boundary

### 15.1 Civilizational OS as aspiration, not proven status

MathGov may be described as a proposed ethical operating system or civilizational operating-system candidate when speaking about its design ambition. That language should remain subordinate to the evidence. The current claimable status is more modest: the MathGov framework is a governance-and-alignment research program with RippleLogic as its ethical decision operating system with a partially calculable ethical decision architecture, an assurance wrapper, and an agent governance companion.

### 15.2 Protection is not participation

SGP is clearest when it separates moral patienthood from governance participation. A being can deserve protection even if it cannot deliberate, communicate, vote, govern, or carry stewardship responsibility. This is especially important for animals, infants, cognitively atypical humans, non-linguistic minds, and future artificial or hybrid systems.

### 15.3 Evidence labels for ordinary users

When MathGov displays a number or score, users should be able to see whether the value was measured, estimated, assumed, expert-judged, contested, unvalidated, not observed, or only a sensitivity placeholder. This reduces false precision and helps people understand what the system knows, what it guesses, and what it has not yet proven.

### 15.4 Compact profiles

For practical adoption, MathGov should be offered through compact profiles: policy review, procurement review, AI deployment review, university AI policy review, and family/community decision support. These profiles can simplify the process, but they cannot waive rights floors, catastrophic-risk checks, containment, refusal, or public accountability.

## Section 16. Historical Lineage: Why v10.8 Focused on Integration, Not Grand Claims

The next meaningful improvement is not a larger claim about civilization. It is better integration: making sure evidence labels, SGP protection boundaries, compact profiles, workbook sanity checks, and release verification all point to the same disciplined claim boundary.

A framework can become more trustworthy by saying less, checking more, and refusing false certainty. The historical v10.8 line should be read in that spirit: a clearer, more integrated release foundation, not a claim of empirical completion.

The practical path remains: reference calculator, schemas, test vectors, replayable PCC examples, independent reconstruction, inter-rater reliability, and shadow-mode pilots.

## Section 17. What the Core Foundation Does and Does Not Include

The core foundation explains the decision architecture, claim boundaries, and public-facing logic of MathGov and RippleLogic. It does not bundle all technical tooling by itself. Schemas, calculators, validators, machine-readable test vectors, evidence-hash registries, and replayable PCC examples belong in companion tooling releases unless explicitly included in the release manifest.

The practical path is staged: v10.8 was the clean core foundation; the current v13.0 line preserves that foundation, carries forward Term Discrimination, and adds Physical Admissibility hardening, while the tooling line develops replayability; later empirical releases must test inter-rater reliability, reconstructability, and shadow-mode pilots before stronger operational-readiness claims are made.

Candidate-readiness interpretation. This primer remains an informative explanatory companion under candidate review. It should not be cited as a normative source of conformance, empirical validation, ProofPack readiness, machine-verifiable ecosystem completeness, or Tier 4 readiness.

### Glossary Binding

All technical terms used in this primer (Rights Floor/RF, NCRC, TRC, CSV, RLS, UCI, HOI, PCC, NCAR, HDW, PLSS, SGP, MPS, FPP, GPR, SPR, ICP, legacy v7 SG\_norm/SG\_patient\_norm fields, rights-admissible, admissible, selectable, selected, Refusal, ProjectionPreRegistration, ComputationalClosure, CategoryCollapseSelfTest, RealityGroundingRecord, RealityGroundingStatus, CMIU, UBR, UBE, MNA, Bridge Premises, Stewardship, Qualification Continuity Check, PCC.ActionQualificationRecord, OutcomeObservation, RequalificationTrigger, ripple.md, Aligners) are defined canonically in RippleLogic Canon v13.0 Appendix L and SGP v8.8. This primer defers to those definitions and adds none of its own.

### Version History

The retained v4.4 release entry appears in Section 23; the entries below preserve earlier v1.x-v2.x lineage. Current component editions and source authority are stated in Appendix RELEASE and the exact build manifest.

v2.6 (Informative). Synchronized to Core Release 2026.07 and companion artifacts: RippleLogic Canon v10.8 / SGP v5.6 / ripple.md v4.0 / Agent System v10.8 / Aligners Sheet v4.0. Added the Reality Grounding and Lean Cascade Patch, Category Grounding / Primitive Audit explanation, structure-transition-consequence vocabulary, measurement-maturity and deployment-context clarification, HDW capture, PLSS escalation, UCI readiness, kernel posture, biological SGP maturity, and release-boundary cleanup.

v2.1 (Informative). Re-pinned to RippleLogic Canon v10.1 / SGP v4.9 / ripple.md v3.4 / Agent System v10.1 / Aligners v3.4; added transition-admissibility and evidence-humility explanation. No normative claims added; companion-to-canon role unchanged.

v2.0 (Informative). Promoted from v1.5 to v2.0 to reflect added Section 13 (What v10.0 Adds: Reality-Coupled Methodology) and re-pinning to RippleLogic Canon v10.0 / SGP v4.8 / ripple.md v3.0 / Agent System v10.0 / Aligners v3.0. No normative claims added; companion-to-canon role unchanged.

v1.5 (Informative). Added a plain-language Qualification Continuity section, updated release-line pins to v10.0 / SGP v4.8 / ripple.md v3.0, and clarified that stability is a transition diagnostic subordinate to rights, tail-risk, containment, and legitimate authority.

v1.4 (Informative). Strengthened the reader guide, added explicit layer discipline, softened union language to avoid final-ontology overclaiming, clarified Refusal in plain language, added a simple low-stakes reading example, expanded glossary binding for v9.6.5 hardening terms, and preserved the v9.6.5 / SGP v4.7.2 / ripple.md v2.3.5 companion line.

v1.2.1 (Informative). Final-fix release. Corrected Section 10 subsection numbering from 9.x to 10.x (carry-over from v1.2.0 renumbering). Replaced docx-js auto-numbering with manually-prefixed paragraphs in Section 6 and Section 11 to fix a rendering issue where some viewers reset every list item to "1." Softened the intelligence-and-isolation passage in Section 4 to avoid unnecessary objections about AI embodiment and absolute scientific claims.

v1.2.0 (Informative). Added Section 4, Union as a Reality Principle. This section unpacks the meaning of UBR for educated general readers, distinguishes union from sameness, names the asymmetric propagation of harm and care, and explicitly rules out three common misreadings: cosmic merger, required opposition (false balance), and collectivist dissolution. Sections 5 through 12 are the v1.0.0 sections renumbered.

v1.0.0 (Informative). Initial Foundations Primer companion to RippleLogic Canon v9.6.5. Established twelve-primitive, eleven-assumption structure with explicit boundary statement, cascade naming, decision-state vocabulary, critical separations, SGP plateau and scalar binding sections, and eleven-step operating sequence including PCC and NCAR.

## Teaching Update: Score, Robustness, and Refusal

RippleLogic first qualifies and then ranks: Reality → Rights → Ruin → Structural viability → Residual ripples. A large downstream score never repairs an upstream failure.

The RLS number is a conditional residual ranking under a declared evidence and value constitution. It is not the amount of welfare in reality, the ethically best option, or permission to execute.

One real effect is counted once. Changing labels, scope views, time slices, reference classes, baselines, confidence, masks or weights cannot silently manufacture a winner. Materially triggered alternatives are tested and disclosed. A unique RLS selection must distinguish the same leader from every selectable rival under every required variant. Testing only the second-highest point score is insufficient.

When the evidence can improve, RippleLogic says EVIDENCE\_UNDERDETERMINED. When legitimate values disagree after the evidence is sufficiently closed, it says NORMATIVE\_UNDERDETERMINED and routes the remaining choice to accountable authority.

## Appendix A: Position, Precedence, and Release Package

This appendix preserves the primer's informative status, pinning, reader guide, and release-package orientation moved out of the opening pages in v2.6. It does not create normative force and does not alter the Canon.

## Section 0. Position and Precedence

This document is an Informative companion to MathGov Core Release 2026.09, governed by RippleLogic Canon v13.0. Its purpose is to make the foundations of MathGov, RippleLogic, and the Sentience Gradient Protocol (SGP) understandable to educated human readers in plain language, without weakening or replacing the normative content of the Canon.

Normative status: This document is Informative. It carries no normative force. Where any sentence in this primer appears to conflict with the canonical specifications, the canonical specifications govern in the order set out in Canon Section 0.1.

Pinning: This primer is pinned to RippleLogic Canon v13.0, SGP v8.8, ripple.md Standard v5.8, WDBIP v1.9, RippleLogic Agent System v13.0, Aligners Sheet v5.9, Cascade Standard v2.9, CSV Gate Standard v2.7, and RLS Validation Protocol v2.9. Any successor release line should re-pin this primer or supersede it explicitly.

Audience: Policy practitioners, governance reviewers, legal staff, public-sector technologists, and educated general readers who need to understand what MathGov claims, what it does not claim, and how its decision architecture works in operational terms.

This primer is explanatory only and does not itself create conformance claims; release identity, claim scope, and normative force remain governed by the pinned canonical artifacts and release manifest.

## Reader Guide

If you are new to MathGov, read Sections 1, 2, 4, 7, 8, and 12 first. These sections explain the purpose, boundaries, union logic, cascade, decision states, and governing posture of the framework in plain language.

If you are reviewing the system for technical or institutional use, read Sections 3, 3A, 3B, 5, 9, 10, 11, and 13 next. These sections identify the declared primitives, assumptions, critical separations, SGP interface, and operating sequence.

If a term in this primer appears to conflict with the canon, the canon controls. This primer is an explanatory bridge, not the source of normative force.

How to read the release package

The Canon is the governing specification. SGP is the moral-status and protection protocol. ripple.md is the portable wrapper and assurance contract. The Agent System is the deployment-control architecture. The Aligners Sheet is a worked-run exemplar and training companion, not a validator. The Methodology Summary explains what changed. The Release Claims and Non-Claims Statement explains what is and is not being claimed publicly.

A simple rule: the core tells you what the system is; the companion files tell you how to understand, test, audit, and eventually validate it.

## Section 18. Historical v2.6 Update: Computable Is Not Automatically Admissible

### 18.1 Plain-language rule

The carried-forward v10.8 line added one important clarity rule: a number can exist without being allowed to decide. A spreadsheet can calculate a score for an option that already failed rights protection. An AI can produce a confident recommendation that lacks authority. A sentience score can be estimated while still being too immature for a strong claim. In those cases, the output is computable but inadmissible.

### 18.2 Why this matters

This protects the framework from false precision. It prevents dashboards, agents, and reports from laundering weak evidence or failed gates into strong decisions. MathGov remains calculable, but calculation remains subordinate to rights, evidence, containment, lawful authority, and claim boundaries.

### 18.3 Adaptive reserve

The carried-forward v10.8 line also added an Adaptive Reserve diagnostic. Some systems fail because they are too chaotic; others fail because they are too tightly optimized and have no slack left. A decision can look efficient while making a family, institution, ecosystem, or society brittle. Adaptive Reserve asks whether enough recovery capacity, redundancy, attention, and reversibility remain after the decision.

## 18.4 v2.6 release delta

The Primer v2.6 explains the Canon v10.8 / SGP v5.6 layer-discipline update in plain language, including Canon Appendix AE. It remains informative only and does not replace the Canon, SGP, ripple.md, Agent System, or Aligners Sheet.

## Section 19. Historical v2.6.1 Update: Framework Verdict Is Not Authority Selection

### 19.1 Plain-language rule

MathGov can say what is blocked, what remains admissible, what is non-decisive, and what requires review. When human, institutional, democratic, expert, emergency, or operator judgment chooses among admissible non-decisive options, that choice must be labeled as an authority selection, not as a unique mathematical verdict.

### 19.2 Why this matters

This protects MathGov from false certainty. A dashboard may compute scores and a governance process may choose a path, but the system must not pretend that a non-decisive score created a unique framework-determined answer. Refusal is not failure; it is the disciplined boundary between evidence-backed decision support and overclaim.

### 19.3 Coverage and projection discipline

Required unknowns must be named. Missing evidence cannot silently become zero. High-stakes runs should state whether scenarios, weights, thresholds, uncertainty assumptions, stakeholder maps, and refusal conditions were declared before evaluation. A changed projection or changed assumption creates a new run boundary.

### 19.4 v2.6.1 release delta

The Primer v2.6.1 explains the MathGov-native hardening layer: decision-verdict semantics, authority-selection separation, projection-freeze visibility, coverage-status discipline, and falsification-class reasoning. It remains informative only and carries no normative force over the Canon or companion specifications.

## Section 20. Historical v2.6 Update Summary: Reality Grounding and Lean Cascade

The v2.6 primer teaches the five-level method: Reality Grounding -> Rights Floor -> TRC -> CSV -> RLS.

Public mantra: Ground -> Protect -> Bound -> Contain -> Score.

UCI/HOI are explained as structural sensors used inside Containment when material and only as special-case tie-break or monitoring tools when RLS is tied, close, or non-decisive.

## Section 21. Category Grounding Patch Summary

The Category Grounding patch clarifies that Category Grounding, Primitive Audit, Reference Category Review, Definition Falsification Test, and Category-Reality Correspondence Check belong to one group inside Reality Grounding when material. They are not separate headline gates. They help MathGov avoid optimizing inherited assumptions or building governance layers on top of unexamined definitions.

Public mantra update: Ground the category, ground the surface, protect the floor, bound ruin, contain structural damage, then score only what remains admissible. The short five-level public method remains Ground -> Protect -> Bound -> Contain -> Score.

## Section 22. Term Discrimination: The Words Must Hold

MathGov now makes one hidden discipline explicit: before a system governs with a word, it must know what the word means, what it excludes, what separates it from neighboring words, and when the word is too weak for the claim being made.

This is Term Discrimination. It belongs inside Category Grounding, which belongs inside Reality Grounding. It is not a new sixth step. The public sequence remains: Reality Grounding -> Rights Floor -> TRC -> CSV -> RLS.

A material term should answer five questions:

1. What does this term mean in this run?
2. What does it exclude?
3. What discriminator separates it from nearby terms?
4. Where does it sit in the dependency order of the decision?
5. What condition forces narrowing, escalation, monitoring-only use, sensitivity-only use, or refusal?

### 22.1 Do-not-collapse matrix

| Do not collapse                                     | Plain meaning                                                                                                                                                         |
|-----------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Reality is not evidence                             | Evidence is contact with a surface of reality, not the whole territory.                                                                                               |
| Evidence is not interpretation                      | Evidence must still be interpreted under uncertainty and values.                                                                                                      |
| A category is not a metric                          | A metric can only measure what its category correctly defines.                                                                                                        |
| A score is not truth                                | RLS ranks remaining options; it does not prove reality.                                                                                                               |
| Compliance is not correctness                       | A system can follow rules that are incomplete or wrong.                                                                                                               |
| Certification is not validity                       | Certification records evidence; it does not manufacture truth.                                                                                                        |
| Permission is not admissibility                     | Approval does not mean the option passed MathGov gates.                                                                                                               |
| Monitoring is not control                           | Watching a system does not mean preventing harm.                                                                                                                      |
| Safety is not stability                             | A safety claim needs bounded conditions; stability needs structural integrity.                                                                                        |
| Stability is not justice                            | A stable system can still violate rights.                                                                                                                             |
| <b>Automation alone does not establish autonomy</b> | A bounded system may be operationally autonomous through coordinated memory, planning, feedback, tools, and controls, but the claim must be defined and demonstrated. |
| Protection is not authority                         | Moral protection does not grant governance                                                                                                                            |

| Do not collapse             | Plain meaning                                                                    |
|-----------------------------|----------------------------------------------------------------------------------|
|                             | power.                                                                           |
| Selected is not executable  | A selected option still needs authority and execution controls.                  |
| Capability is not authority | Being able to act does not make action legitimate.                               |
| Feasibility is not safety   | A physically possible transition may still be unsafe, unlawful, or inadmissible. |
| Inference is not warrant    | A plausible candidate is not the justification required to execute it.           |

Plain-language rule. Bad words create bad measurements. Bad measurements create bad optimization. Bad optimization creates real harm. MathGov therefore grounds the words before it trusts the score.

For material claims about automation, autonomy, agency, reasoning, understanding, self-improvement, intelligence, or intent, do not force the system onto one ladder. State the operational definition, observable criteria, system composition, evidence, validity conditions, alternative explanations, falsifier, maturity stage, permitted wording, and explicit nonclaims. Also record the candidate-generation method, source of constraints, source and change authority for objectives, admissibility warrant, execution interface, execution authority, and revocation or safe-state path separately. A system may combine learned inference, memory, planning, feedback, tools, deterministic checking, optimization, formal methods, and human judgment into a bounded system-level capability. What matters is what the pinned system demonstrates, what remains unexplained, which component proposes, which component justifies, which authority permits, and which control can stop the action.

Also keep the domain of the claim separate from the domain of its warrant. Physical evidence can support physical feasibility within stated limits; it does not by itself establish legal permission, ethical acceptability, or governance legitimacy. A cross-domain conclusion needs an explicit bridge, source of authority, and residual uncertainty.

## 22.2 Public examples

A school that defines learning only as test-score improvement may optimize the dashboard while weakening real understanding. A company that defines productivity only as output per hour may improve reports while hollowing trust, health, and retention. An AI lab that defines safety as monitoring may watch risks it has not actually controlled. A regulator that treats certification as validity may approve a system whose category, evidence, or reference structure was never strong enough for the claim.

Term Discrimination is the shield against this drift. It does not make MathGov more complicated for its own sake. It is intended to make the framework harder to game, easier to audit, and more explicit about what its words can support; practical benefits require empirical testing.

## Section 23. Release Lineage

### v4.4 Retained Release Entry (Lineage)

LINEAGE (historical). v4.4 synchronized the Primer with Canon v12.6, SGP v8.6, ripple.md v5.6, WDBIP v1.7, Agent System v12.6, Aligners Sheet v5.7, and run-record v4.1.1 schema contract. It preserves the public two-phase and five-stage explanation while adding only bounded implementation guidance: consequence-tempo assurance, responsibility continuity, welfare-boundary discipline, and action-bound qualification continuity. These do not add public stages, create empirical validation, or grant execution authority.

## v4.0 Retained Release Entry (Lineage)

v4.0 (Informative). Synchronizes the Primer with MathGov Core v12.2 / Canon v12.2 / SGP v8.1-line / ripple.md v5.0-line / Agent System v12.0.1 / Aligners Sheet v5.2.1. It teaches the two-phase method: first qualify options through RG -> RF -> TRC -> CSV, then rank survivors through RLS. It carries forward Term Discrimination and Semantic Stability while adding the Physical Admissibility boundary: governance permission is not physical safety proof, generated candidates are not admissibility warrants, and PC-AEP belongs inside RG and CSV when material.

## Historical Primer Lineage Note

Earlier v1.x-v4.0 entries below are retained as lineage notes only. They do not control current v4.7 interpretation.

The retained Primer v4.0 entry records the Canon v12.2 / SGP v8.1-line / ripple.md v5.0-line / Agent v12.0-line / Aligners v5.2-line synchronization at that time. The major v12.0 addition is the Physical Admissibility Release, carrying forward the Term Discrimination and Semantic Stability Patch. It adds explicit CSV treatment for Containment and Structural Viability while carrying forward the public discipline of naming the discriminator, dependency position, non-collapse risks, and refusal condition for material terms.

Public mantra update: Ground the words, ground the surface, protect the floor, bound ruin, contain structural damage, then score only what remains admissible. The short five-level public method remains Ground -> Protect -> Bound -> Contain -> Score.

## Computable, realizable, and ethical are not the same

A modern AI system can generate far more than reality can carry. RippleLogic therefore refuses the shortcut from possibility to selection. A generated option may be computable, fluent, executable by a tool, or attractive in a dashboard and still fail the cascade.

The public rule is: **Computable is not selectable. Generated is not grounded. Executable is not ethical. Realizable is necessary but not sufficient.**

This is why the cascade begins with Reality Grounding and ends with RLS only after Rights Floor, TRC, and CSV have formed the selectable set.

HISTORICAL LINEAGE - NOT CURRENT SPEC. The following passages preserve prior component pins for provenance only. Current pins are stated in Appendix RELEASE and the release manifest.

Historical lineage note: release-history sections are informative and are not current governing specifications. Current v4.7 pins are controlled by the release manifest.

## Methodological integrity primer note

The v12.0 system adds Methodological Falsifiability and Dependency Integrity as a claim-discipline inside Reality Grounding and downstream standards. Agents and reviewers must not treat plausible output, coherence, approval, or post-hoc tuning as support for stronger claims than the evidence, dependencies, alternatives, and revision triggers permit.

END HISTORICAL LINEAGE. The framework relationship map and following current provisions are governed by the Canon and the current component matrix in Appendix RELEASE.

## Framework relationship map

| Layer       | Role                                   |
|-------------|----------------------------------------|
| MathGov     | Umbrella ethical governance framework. |
| RippleLogic | Decision architecture inside MathGov.  |

| Layer          | Role                                               |
|----------------|----------------------------------------------------|
| SGP            | Companion moral-status and protection protocol.    |
| ripple.md      | Portable wrapper and Decision Note standard.       |
| Agent System   | Runtime control architecture for AI/hybrid agents. |
| Aligners Sheet | Worked-run exemplar, not validator.                |

### Three use profiles: easy when simple, complete when necessary

MathGov uses one architecture with proportional documentation. **Quick / Tier 1** is for low-stakes, reversible personal, family, and classroom choices and may use a short qualitative record, including a required qualitative ruin-trigger screen; triggered ruin concerns require escalation and governed TRC assessment. **Standard / Tier 2** adds evidence, stakeholder, parameter, gate, uncertainty, control, and audit records for organizational or community decisions. **Audit / Tier 3** adds exact hashes, independent review, scenario and subgroup governance, and replay for public, contested, irreversible, or high-stakes decisions.

The controlling reading specification is

Core\_15/MATHGOV\_REPRODUCIBILITY\_AND\_USE\_STANDARD\_v1.7.docx. A lower profile cannot be used to avoid a material rights, ruin, structural, evidence, or authority requirement. The compact twelve-rule kernel and the requirements for a run schema, validator, and worked replay are specified in the candidate MathGov Reproducibility and Use Standard v1.7. The baseline implementation reference is docs/reproducibility/MATHGOV\_REPRODUCIBILITY\_AND\_USE\_STANDARD\_v1.5.md; its candidate source parity remains to be established.

## APPENDIX RELEASE: Identity, Source Authority and Revision Record

Framework release: MathGov/RippleLogic v13.0. Component: RippleLogic Foundations Primer v4.7. 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-13.0-20260926-RELEASE-I. Prepared 26 September 2026 from the supplied Release H publication package. Only reproduced errata and necessary source/evidence synchronization were patched; component editions are unchanged. The manifest identifies each component's exact bytes and source build. Live publication is not asserted.

Component/build identity. The manifest distinguishes this correction build from unchanged Release G and H masters. Aligners Sheet v5.9 keeps its numerical inputs, formulas, cached results and worked verdict; text and print-area errata are itemized in Verification/Final\_I/Exact\_Workbook\_Changes.json. A filename or component edition alone is not exact file identity.

Verification boundary: Reports/Release\_I\_Verification.md and Verification/Final\_I/ identify the checks performed on this package and their limits. Prior-build receipts do not certify changed bytes. Cache-independent replay and spreadsheet-application recalculation 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            |
| ripple.md Standard          | v5.8            |
| RippleLogic Agent System    | v13.0           |

| Component                                                       | Current edition |
|-----------------------------------------------------------------|-----------------|
|                                                                 |                 |
| 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/Release\_I\_Adjudication.md; exact edits: Verification/Final\_I/Exact\_Document\_Changes.json and Exact\_Workbook\_Changes.json. Navigation and preservation records are separate. The untouched prior publication and original 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)

Version 4.3 (Informative lineage; not the current release)

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.

### v4.6-rc2 Historical Release Integration

HISTORICAL (NON-CONTROLLING): Release: MathGov Core Release 2026.09 — RippleLogic Canon v12.8-rc2 / SGP v8.7-rc2 — Rights Non-Dilution, Representation Integrity, and Robust Selection Release

Historical companion pins: Canon v12.8-rc2; SGP v8.7-rc2; ripple.md v5.7-rc2; Agent System v12.7-rc2; CSV v2.6-rc2; Cascade v2.8-rc2; Reproducibility v1.6-rc2; WDBIP v1.8-rc2; RLS Validation v2.8-rc2; Primer v4.6-rc2; Public Introduction v12.8-rc2; PC-AEP/MFDI/Source-Coupling v2.5-rc2; Aligners Sheet v5.8-rc1.

| Release control | Historical value                                                                                        |
|-----------------|---------------------------------------------------------------------------------------------------------|
| Component       | RippleLogic Foundations Primer v4.6-rc2                                                                 |
| Release         | HISTORICAL (NON-CONTROLLING): MathGov Core Release 2026.09 — RippleLogic Canon v12.8-rc2 / SGP v8.7-rc2 |

| Release control      | Historical value                                                                                                                                                                                                                                       |
|----------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Architecture         | RG → RF/NCRC → TRC → CSV → RLS                                                                                                                                                                                                                         |
| Role                 | Informative human doorway to the current Core; cannot override the Canon.                                                                                                                                                                              |
| Claim boundary       | Release candidate for Tier 1–3 research; source/mirror parity and implementation conformance remain separately reviewable; not empirical validation, legal authority, physical-safety certification, deployment authorization, Tier 4, or moral truth. |
| Source/render parity | Versioned semantic source and DOCX/PDF mirrors must agree. Filename, internal version, active pins, manifest, and hashes must agree; mismatch is release-integrity failure.                                                                            |

## Historical Release Integration (v4.6-rc2)

**HISTORICAL (NON-CONTROLLING):** This candidate advances RippleLogic Foundations Primer from the audited v4.5 baseline to v4.6-rc2 and binds it to MathGov Core Release 2026.09 — RippleLogic Canon v12.8-rc2 / SGP v8.7-rc2. Historical references remain lineage only; unmarked current-release references use the current component matrix in Appendix RELEASE.

### v4.1.1 Patch Note (Lineage)

v4.1.1 synchronizes the public explanatory surface to the v12.4.1 package. It remains informative and adds no rule.

Companion to Canon v12.8-rc2 / SGP v8.7-rc2 / ripple.md v5.7-rc2 / WDBIP v1.8-rc2 / Agent System v12.7-rc2 / Primer v4.6-rc2 / Aligners Sheet v5.8-rc1

### v4.2 implementation note (Lineage; retained for explanation)

The public five-stage method remains unchanged. In the v12.7 lineage, two implementation disciplines were carried beneath the doorway: controls had to be able to act before unacceptable adverse escalation, and distributed work had to preserve answerability. WDBIP v1.7 was the then-current welfare-boundary companion; the present candidate is WDBIP v1.8-rc2. These are deeper implementation aids, not extra public steps.

## Release alignment (v12.8-rc2)

**HISTORICAL (NON-CONTROLLING):** This companion document is synchronized with MathGov Core Release 2026.09 — RippleLogic Canon v12.8-rc2 / SGP v8.7-rc2 — Rights Non-Dilution, Representation Integrity, and Robust Selection Release. The public cascade is RG → RF → TRC → CSV → RLS; the formal shorthand is RG/RSG → RF/NCRC → TRC → CSV → RLS. RF is the public Rights Floor layer and NCRC is its formal non-compensatory predicate. CSV means Containment and Structural Viability. Detailed cascade definitions are maintained in the Canon, the RippleLogic Cascade Standard, and the CSV Gate Standard; this file applies those invariants without repeating the full release note.

**LINEAGE (historical).** Version 4.3 was a synchronized companion edition for Canon / SGP / ripple.md / WDBIP / Agent System / Primer / Aligners Sheet. Package line: MathGov Core Release 2026.09 — RippleLogic Canon / SGP — Rights Non-Dilution, Representation Integrity, and Robust Selection Release. It

carries forward the Physical Admissibility boundary and the Term Discrimination and Semantic Stability discipline.