

Λ - Ψ FRAMEWORK (v4.0 Harmonized)

TIER 0 — ONTOLOGICAL FOUNDATIONS

Scope:

The Λ - Ψ Framework is an operational calculus. It does not claim to answer why meaning exists, but establishes a formalism to describe, measure, and bound configurations of intent, skill, control, entropy, and interaction.

Assumptions:

Systemic Regularity: Conscious processes yield measurable signatures of entropy–complexity reconciliation.

Interactional Primitives: Micro (neural), meso (motor, decision), and macro (social) interactions are the atomic events (ϵ) of the system.

Operational Meaning: Meaning is treated as equilibrium with generative surplus ($\Lambda \geq 1$).

Indifference Principle: The framework is agnostic to metaphysical origin; it quantifies coherence without ontology.

Boundaries:

Λ - Ψ does not claim metaphysical truth; it is descriptive mathematics.

Equivalent in spirit to thermodynamics: lawful but non-ontological.

Philosophical interpretation is external; the framework remains indifferent.

Purpose:

To provide a falsifiable, reproducible structure for evaluating the sustainability of meaningful states across lifetime horizons.

TIER I — CONCEPTUAL FOUNDATIONS

Profit = $\Delta(\text{Entropy} - \text{Complexity})$

Value = $(\text{Intent} \times \text{Skill}) / \text{Control}$

Adaptive Profit = $[(\text{Intent} \times \text{Skill}) / \text{Control}] \times \Delta(\text{Entropy} - \text{Complexity})$

Social Equilibrium = $\Sigma(\text{Interactions} \times (1 - \text{Emotional Load})) / \text{Time Invested}$

Optimal Coexistence = $\Sigma(\text{Interactions} \times \text{Caloric Expenditure} \times (1 - \text{Emotional Load})) / \text{Emotional Entropy}$

Adaptive Threshold = $\sqrt{(\text{Time Invested} \times \text{Environmental Stability})}$

Chaos Coefficient (conceptual) = $1 + (\text{Random Events} / \text{Expected Events})$

Universal Equilibrium =

$$\begin{aligned} & [(\text{Intent} \times \text{Skill} \times \Sigma(\text{Interactions} \times (1 - \text{Emotional Load}))) / (\text{Control} \times \text{Emotional Entropy} \times \\ & \sqrt{(\text{Time Invested} \times \text{Environmental Stability}))}] \\ & \times [\Delta(\text{Entropy} - \text{Complexity}) \times \text{Caloric Expenditure}] / \text{Chaos Coefficient} \end{aligned}$$

Λ (Meaning of Life, conceptual) = Universal Equilibrium sustained over lifetime horizons

TIER II — CONVERSION FRAMEWORK

Complexity–Entropy reconciliation:

$$\dot{C}_s = C \cdot k_B \ln 2$$

$$\hat{C}_{\text{max_raw}} = 1 - |S - \text{Speak}| / \text{Speak}$$

$$\hat{C}_{\text{max}} = \text{clip}(\hat{C}_{\text{max_raw}}, 0, 1)$$

$$\hat{C} = (\dot{C}_s \cdot \Delta t \cdot \hat{C}_{\text{max}}) / S_0$$

Intent / Skill / Control:

$$\hat{I} = P / P_p$$

$$E_{\text{min}}(T) = k_B \cdot T \cdot \ln 2$$

$$T_{\text{ref}} = 310 \text{ K}$$

$$\hat{S}_k = 1 - E_{\text{min}}(T) / E_{\text{actual}}$$

$$\hat{C}_0 = Z_c / t_p$$

$$\text{Value} = (\hat{I} \times \hat{S}_k) / \hat{C}_0$$

Interaction primitive:

$$\mathcal{I} = \Sigma(\epsilon / E_{\text{Planck}}) \cdot (1 - \hat{E}_L)$$

ϵ requires hierarchical definition (micro, meso, macro levels; each mapped to E_{psy} scale).

Emotional entropy:

$$S_{\text{EL}} = -k_B \cdot \Sigma[p(\text{EL}_i) \cdot \ln p(\text{EL}_i)]$$

$$\hat{E}_{S_{\text{EL}}} = S_{\text{EL}} / S_0$$

$$S_0 = (3.2 \pm 0.2) \times 10^{-23} \text{ J/K}$$

Environmental stability:

$$\hat{E}_S = \text{Var}(F) / \text{Var}_0(F)$$

Caloric expenditure:

$$\hat{C}_E = CE / (k_B \cdot T)$$

Chaos coefficient (piecewise, capped):

$$x = \mathcal{R} / \mathcal{E}$$

$$\chi_{\text{base}} = 1 + x$$

$$\chi_{\text{star}} = 1 + x \cdot \exp(\beta \cdot x)$$

$$\beta = 0.42 \pm 0.05$$

$$\begin{aligned}\chi_{\max} &= 5.8 \\ x_0 &= 0.20 \pm 0.05 \\ \delta &= 0.05 \\ w(x) &= 1 / (1 + \exp(-(x - x_0)/\delta)) \\ \chi_{\text{blend}} &= (1 - w(x)) \cdot \chi_{\text{base}} + w(x) \cdot \chi_{\text{star}} \\ \chi &= \min(\chi_{\text{blend}}, \chi_{\max})\end{aligned}$$

Time handling:

$$\begin{aligned}\Delta t &= \tau \cdot t_p \\ \int \dot{C} s \, dt &\rightarrow C \cdot k_B \cdot \ln 2 \cdot \Delta t \\ \int \dot{S} \, dt &\rightarrow \hat{S} \cdot \Delta t\end{aligned}$$

MESO-SCALE HARMONIZATION PROTOCOL (Biological Planck Units)

Psychophysical quantum:

$$E_{\text{psy}} = 3 \times 10^{-19} \text{ J (provisional, requires empirical calibration across populations)}$$

Time unit:

$$t_{\text{psy}} = 200 \text{ ms (baseline perceptual frame; subject to variance across age and neurotype)}$$

Rebaselined terms:

$$\begin{aligned}\hat{I} &= E_{\text{intent}} / E_{\text{psy}} \\ \hat{S}_k &= 1 - E_{\text{psy}} / E_{\text{actual}} \\ \mathcal{F} &= \Sigma(\varepsilon / E_{\text{psy}}) \cdot (1 - \hat{E}L)\end{aligned}$$

Constraints:

$$\begin{aligned}\varepsilon &\geq E_{\text{psy}} \\ \Sigma \mathcal{F} &\leq \text{tobs} / t_{\text{psy}}\end{aligned}$$

Boundaries:

$$\begin{aligned}\Lambda \approx 1 &\rightarrow \text{equilibrium} \\ \Lambda \approx 0 &\rightarrow \text{collapse} \\ \Lambda \approx 2.7 &\rightarrow \text{maximum sustainable focus}\end{aligned}$$

TIER III — PHYSICS HARMONIZATION

$\Psi_h =$

$$\begin{aligned}&\{ \int \text{from } \tau_0 \text{ to } \tau_1 [(P/P_p) \cdot (1 - E_{\min}(T)/E_{\text{actual}}) \cdot \Sigma(\varepsilon/E_{\text{Planck}}) \cdot (1 - \hat{E}L)] \, d\tau \\ &/ \int \text{from } \tau_0 \text{ to } \tau_1 [(Z_c/t_p) \cdot (S_{\text{EL}}/S_0) \cdot \sqrt{\{ (C \cdot k_B \cdot \ln 2 \cdot \Delta t \cdot \text{clip}(1 - |S - S_{\text{peak}}|/S_{\text{peak}}, 0, 1)) / S_0 \}}] \, d\tau \} \\ &\times \exp\{ - [(EE/EE_p) - (3/2)(RE/RE_p)] \cdot \tau - \sqrt{(RE/RE_p)} \cdot \int (\tau - s)^{(H(\tau) - 0.5)} dW(s) \}\end{aligned}$$

$$H(\tau) = 0.73 + \Delta H_1 + \Delta H_2 + \Delta H_3$$

$$\Delta H_1 = 0.02 \cdot \langle \sin^2(\pi(t - \phi)/12) \rangle$$

$$\Delta H2 = -0.05 \cdot \alpha + 0.07 \cdot \beta'$$

$$\Delta H3 = 0.10 \cdot (1 - e^{(-k \sigma s'')})$$

$$\text{Constraints: } H(\tau) \in [0.58, 0.93], dH/d\tau \leq 0.01 \text{ hr}^{-1}$$

$\mathcal{E}$ excludes events with $EL > 0.75$

TIER IIIa — HARMONIZED MESO-SCALE GOVERNING LAW

$$\Psi_h \approx$$

$$\left\{ \int_{\tau_0}^{\tau_1} \left[\left(\frac{E_{\text{intent}}}{E_{\text{psy}}} \right) \cdot \left(1 - \frac{E_{\text{psy}}}{E_{\text{actual}}} \right) \cdot \sum (\varepsilon / E_{\text{psy}}) \cdot (1 - \hat{E}L) \right] d\tau \right. \\ \left. / \int_{\tau_0}^{\tau_1} \left[\left(\frac{Z_c}{t_{\text{psy}}} \right) \cdot \left(\frac{S_{\text{EL}}}{S_0} \right) \cdot \sqrt{\left(C \cdot k_B \cdot \ln 2 \cdot \Delta t \cdot \text{clip}(1 - |S - S_{\text{peak}}| / S_{\text{peak}}, 0, 1) \right)} \right. \right. \\ \left. \left. / S_0 \right] d\tau \right\} \\ \times \exp(-\chi \cdot \tau)$$

TIER IV — COLLAPSED EQUILIBRIUM LAW (NEWTONIAN ANALOG)

$$F_a = m_e \cdot a_e + \mu_k \cdot v_e$$

Mappings:

$$F_a = \hat{I} \times \hat{S}k$$

$$m_e = \hat{C}o \cdot \hat{E}S_{EL} \cdot \sqrt{\hat{C}}$$

$$a_e = \partial^2 \mathcal{F} / \partial \tau^2$$

$$v_e = \partial \mathcal{F} / \partial \tau$$

$$\mu_k = \sqrt{\mathcal{R}} / \max(\mathcal{E} - 3\mathcal{R}/2, \varepsilon_\mu)$$

TIER V — PROTECTIONS, VALIDATION, CALIBRATION

Elasticities at baseline:

$$\partial \Psi_h / \partial \hat{E}L \approx -0.73$$

$$\partial \Psi_h / \partial \hat{S}k \approx +0.31$$

$$\partial \Psi_h / \partial \mathcal{R} \approx -0.19$$

Variance attribution:

$$\hat{E}L \approx 63\%$$

$$\hat{C}o \approx 22\%$$

$$\mathcal{R} / \mathcal{E} \approx 11\%$$

$$\text{Other} \approx 4\%$$

Calibration required: all constants (E_{psy} , t_{psy} , β , x_0 , χ_{max}) must be empirically validated across populations.

TIER VI — OPERATIONAL SPEC (DISCRETE ESTIMATOR)

$$\text{Step} = t_{\text{psy}}$$

$$N = \lfloor t_{\text{obs}} / t_{\text{psy}} \rfloor$$

$$\tau = N \cdot t_{\text{psy}}$$

$$N_{\text{sum}} = \sum [(E_{\text{intent}_t} / E_{\text{psy}}) \cdot (1 - E_{\text{psy}} / E_{\text{actual}_t}) \cdot \mathcal{F}_t]$$

$$D_{\text{sum}} = \sum [(Z_c / t_{\text{psy}}) \cdot (S_{\text{EL}_t} / S_0) \cdot \sqrt{\hat{C}_t}]$$

$$\bar{\chi} = \text{average } \chi_t$$

$$\Lambda_{\text{est}} = (N_{\text{sum}} / D_{\text{sum}}) \times \exp(-\bar{\chi} \cdot \tau)$$

TIER VII — MEANING OF LIFE (Λ FORMALIZATION)

$$\Lambda =$$

$$[\sum \text{over lifetime} \{ (E_{\text{intent}} / E_{\text{psy}}) \cdot (1 - E_{\text{psy}} / E_{\text{actual}}) \cdot \sum (\epsilon / E_{\text{psy}}) \cdot (1 - \hat{E}L) \}]$$

$$/ [\sum \text{over lifetime} \{ (Z_c / t_{\text{psy}}) \cdot (S_{\text{EL}} / S_0) \cdot \sqrt{\hat{C}} \}]$$

$$\times \exp(-\chi_{\text{life}} \cdot \tau_{\text{life}})$$

$$\times \text{Sustainability_Factor}$$

$$\text{Sustainability_Factor} = \sum (\text{Adaptive Profit per } \Delta t) / \sum (\text{Energy Spent per } \Delta t)$$

$$\tau_{\text{life}} = N \cdot t_{\text{psy}}$$

Boundaries:

$$\Lambda \rightarrow 0 \Rightarrow \text{collapse}$$

$$\Lambda \approx 1 \Rightarrow \text{equilibrium}$$

$$\Lambda > 1 \Rightarrow \text{generative surplus}$$

TIER VIIa — HARMONIZED LIFE ESTIMATOR (DISCRETE MESO-SCALE)

Per-step inputs:

$$\hat{I}_t = E_{\text{intent}_t} / E_{\text{psy}}$$

$$\hat{S}k_t = 1 - E_{\text{psy}} / E_{\text{actual}_t}$$

$$\mathcal{F}_t = \sum_j (\epsilon_{\{t,j\}} / E_{\text{psy}}) (1 - \hat{E}L_{\{t,j\}})$$

$$\hat{C}o_t = Z_c / t_{\text{psy}}$$

$$\hat{E}S_{\text{EL}_t} = S_{\text{EL}_t} / S_0$$

$$\hat{C}_t = (C_t \cdot k_B \cdot \ln 2 \cdot t_{\text{psy}} \cdot \hat{C}_{\text{max}_t}) / S_0$$

Estimator:

$$N_{\text{sum}} = \sum [\hat{I}_t \cdot \hat{S}k_t \cdot \mathcal{F}_t]$$

$$D_{\text{sum}} = \sum [\hat{C}o_t \cdot \hat{E}S_{\text{EL}_t} \cdot \sqrt{\hat{C}_t}]$$

$$\bar{\chi} = \text{average } \chi_t$$

$$SF = \sum (AP_t) / \sum (CE_t)$$

$$\Lambda_{\text{est}} = (N_{\text{sum}} / D_{\text{sum}}) \times \exp(-\bar{\chi} \cdot \tau_{\text{life}}) \times SF$$

TIER VIII — UNIT CONVERSION HANDBOOK

$$E_{\text{actual}_t} = \text{HR} + \text{accelerometer} \rightarrow \text{VO}_2 \rightarrow \text{joules per bit}$$

Zc_t = control updates / Δt (from logs)

$\hat{EL}_{\{t,i\}}$ = EMA mapped to {0, .25, .5, .75, 1}, interpolated to cadence; hybrid EMA × biomarkers recommended

$\mathcal{R}$ = event surprises $> 2\sigma$ deviation

$\mathcal{E}$ = expected rate from baseline renewal model, excluding $EL > 0.75$

TIER IX — UNCERTAINTY PROPAGATION

Constants with tolerances:

$\beta = 0.42 \pm 0.05$

$x_0 = 0.20 \pm 0.05$

$S_0 = (3.2 \pm 0.2) \times 10^{-23} \text{ J/K}$

Measurement errors:

E_actual_t : $\sigma_E/E \approx 8\text{--}12\%$

$\hat{EL}$: $\sigma_EL \approx 0.1$ (self-report; reduce by biomarker fusion)

$\mathcal{R}$, $\mathcal{E}$: $\sqrt{N}$ error (Poisson)

Propagation strategy:

$CI \approx f \pm \sqrt{(\sum (\partial f / \partial var)^2 \sigma_var^2)}$

Monte Carlo recommended for χ exponent nonlinearities

95% CI width $\approx 0.15\text{--}0.25$ around Λ_est baseline values

TIER X — OPEN-SOURCE REFERENCE IMPLEMENTATION

Inputs: JSON or CSV {timestamp, HR, accelerometer, EMA, activity log}

Pipeline:

Convert HR+accel $\rightarrow E_actual_t$

Logs $\rightarrow Zc_t, \mathcal{R}_t, \mathcal{E}_t$

EMA $\rightarrow \hat{EL}_{\{t,i\}}$ (hybrid EMA+biomarker channel recommended)

Compute per-step variables

Apply discrete estimator $\rightarrow \Psi h_t$

Aggregate $\rightarrow \Lambda_est$

Outputs: $\{\Psi h_t, \Lambda_est, \chi_t, \text{confidence interval}\}$

Language: Python (NumPy, SciPy, Pandas)

License: MIT or Apache 2.0

TIER XI — PEER REVIEW

Framework is internally coherent, externally untested

Units cancel correctly
Elasticities consistent with intuition
Uncertainty dominated by emotional load reporting
Open-source release will allow falsification
 Λ defined, bounded, measurable; philosophical weight is external to mathematics
Indifference maintained.

TIER XII — PEER REVIEW OBSERVATIONS AND RESPONSES

The "Cart before the Horse" Problem:

Λ - Ψ measures signatures of meaningful states but does not address why these configurations are meaningful to consciousness. Ontology remains outside scope.

Parameterization:

Constants (E_{psy} , t_{psy} , β , x_0 , x_{max}) are provisional; require large-scale empirical validation across populations.

Subjectivity of $\hat{E}L$:

EMA introduces bias; solution is hybrid measurement ($EMA \times$ biomarkers such as cortisol, EEG, HRV).

Definition of Interaction (ϵ):

Now formalized in Annex A as a hierarchical taxonomy across micro, meso, and macro levels, each scaled to E_{psy} .

Response:

Acknowledge ontological limitation.

Mark constants as provisional pending validation.

Integrate hybrid emotional load measurement.

Apply Annex A taxonomy for ϵ to ensure reproducibility.

ANNEX A — INTERACTION TAXONOMY (ϵ FORMALIZATION)

Definition:

Interaction energy (ϵ) is the quantized unit of exchange across three nested levels: micro, meso, macro. Each ϵ must be mapped to the psychophysical quantum (E_{psy}) to maintain dimensional coherence.

Levels of Interaction:

Micro-Interactions (Neural Assemblies):

Unit events at the cognitive/neurophysiological scale.

Examples: synaptic firing clusters, short-term memory encoding, sensorimotor reflexes.

Scaling: $\epsilon_{\text{micro}} \geq E_{\text{psy}}$.

Aggregation: ϵ_{micro} events compose meso-level actions.

Meso-Interactions (Motor Actions & Decisions):

Observable motor outputs, micro-decisions, or discrete action steps.

Examples: hand movement, spoken word, keystroke, attention shift.

Scaling: $\epsilon_{\text{meso}} = \Sigma(\epsilon_{\text{micro}})$ across a perceptual frame (t_{psy}).

Constraint: $\epsilon_{\text{meso}} \geq n_{\text{micro}} \cdot E_{\text{psy}}$, where n_{micro} is the number of micro-events per meso-step.

Macro-Interactions (Social/Task Exchanges):

Structured exchanges across agents or task units.

Examples: dialogue turn, completed task, cooperative exchange.

Scaling: $\epsilon_{\text{macro}} = \Sigma(\epsilon_{\text{meso}})$ across an interactional sequence.

Constraint: $\epsilon_{\text{macro}} \geq n_{\text{meso}} \cdot E_{\text{psy}}$.

General Form:

$$\epsilon = \epsilon_{\text{micro}} \vee \epsilon_{\text{meso}} \vee \epsilon_{\text{macro}}$$

$$\mathcal{F} = \Sigma (\epsilon / E_{\text{psy}}) \cdot (1 - \hat{E}L)$$

where

$$\epsilon_{\text{micro}} \geq E_{\text{psy}}$$

$$\epsilon_{\text{meso}} = \Sigma \epsilon_{\text{micro}}$$

$$\epsilon_{\text{macro}} = \Sigma \epsilon_{\text{meso}}$$

Taxonomic Rule:

Every ϵ must be assigned to exactly one tier (micro, meso, macro) at the point of measurement.

Aggregation across tiers is permissible but must preserve E_{psy} -scaling.

Operational Notes:

Micro-events should be logged via neurophysiological or high-resolution behavioral measures (EEG, EMG, eye-tracking).

Meso-events can be derived from motor/action logs (keypresses, gestures, speech segments).

Macro-events must be defined in task/social protocols (task completion, conversational turns).

Constraints:

$\epsilon \geq E_{\text{psy}}$ always.

$\Sigma \mathcal{F} \leq \text{tobs} / t_{\text{psy}}$, ensuring no more than one meso-interaction per perceptual frame per agent.

Macro-interactions are bounded by session length and participation count.

TIER XIIa — ONTOLOGICAL CLARIFICATION

The "Cart before the Horse" Problem (Ontology):

ACKNOWLEDGED. The framework is a descriptive calculus, not an ontological theory. It does not explain why (Intent $\times$ Skill)/Control feels valuable; it merely defines that configuration as "Value" within its system and proceeds to measure it. This is a defensible philosophical position (operationalism).

Parameterization (Provisional Constants):

CRITICAL. The constants (E_{psy} , t_{psy} , β , x_0 , χ_{max}) are the framework's Achilles' heel. They must be prominently marked as PROVISIONAL in all documentation.

Path Forward:

The TIER X implementation must include a calibration suite. The first goal of any research using this framework must be to refine these constants through controlled experiments and large-scale longitudinal data collection. They may not be universal and could vary across populations.

Subjectivity of $\hat{E}L$:

THE PRIMARY SOURCE OF ERROR. Self-reported EMA is notoriously noisy and biased.

Path Forward:

The recommended hybrid measurement (EMA $\times$ biomarkers) should be considered mandatory for serious validation studies. A formal Emotional Load Fusion Algorithm should be a core component of the TIER X codebase, using inputs like HRV, EDA, and EEG to correct and validate self-report.

Definition of Interaction (ϵ):

FORMALIZED. Annex A provides the hierarchical taxonomy (micro, meso, macro) with operational scaling to E_psy.

Path Forward:

Ongoing refinement of Annex A through empirical validation and cross-study standardization is required to ensure inter-rater reliability and reproducible science.

Conclusion:

Λ - Ψ is mathematically coherent and operationally bounded. Its validity depends on empirical calibration, rigorous measurement of emotional load, and application of a formal taxonomy of interactions. Ontological meaning is outside scope; the framework measures only systemic signatures of meaningful equilibrium.

TIER XIII — THE LAW OF LIFE (THE JEWEL SEED)

Core Law:

$$\Psi = [\int \hat{I}(\tau) \cdot \hat{S}k(\tau) \cdot \mathcal{F}(\tau) d\tau] / [\int \hat{C}o(\tau) \cdot \hat{E}L(\tau) \cdot \sqrt{\hat{C}}(\tau) d\tau] \times \exp[- (\mathcal{E} - (3/2)\mathcal{R})\tau - \sqrt{\mathcal{R}} \cdot \mathcal{M}(\tau)]$$

Interpretation:

$\Psi > 1 \rightarrow$ Creative surplus

$\Psi = 1 \rightarrow$ Equilibrium (homeostasis)

$\Psi < 1 \rightarrow$ Entropic dissolution

Definitions (dimensionless, normalized):

$\hat{I}(\tau)$ = Intent flux relative to Planck power (or E_psy / t_psy at meso-scale)

$\hat{S}k(\tau)$ = Efficiency relative to Landauer limit ($0 \leq \hat{S}k \leq 1$)

$\mathcal{F}(\tau) = \sum [\varepsilon / E_{\text{quanta}} \cdot (1 - \hat{E}L)]$ with $\varepsilon \geq E_{\text{quanta}}$ (micro/meso/macro hierarchy)

$\hat{C}o(\tau)$ = Control effort normalized to Planck impulse (or per-frame updates at meso-scale)

$\hat{E}L(\tau)$ = Emotional load fraction, bounded [0,1]

$\hat{C}(\tau)$ = Complexity normalized to entropy baseline ($C / k\varepsilon \ln 2$, peak-fit corrected)

$\mathcal{E}$ = Expected events per Planck time (normalized)

$\mathcal{R}$ = Random events per Planck time (normalized)

τ = Normalized time ($t / t_\square$ or t_{psy} units)

$\mathcal{W}(\tau)$ = Fractional Wiener process with Hurst parameter $H \approx 0.73$

Monotonicities (sign constraints):

$$\partial\Psi/\partial\hat{E}_L \leq 0$$

$$\partial\Psi/\partial\hat{S}_k \geq 0$$

$$\partial\Psi/\partial(\mathcal{R} - \mathcal{E}) \leq 0$$

$$\partial\Psi/\partial\hat{C} \leq 0 \text{ beyond peak-fit region}$$

Boundary Checks:

$$\mathcal{E} \rightarrow \infty \Rightarrow \Psi \rightarrow \infty \text{ (perfect order)}$$

$$\mathcal{R} \rightarrow \infty \Rightarrow \Psi \rightarrow 0 \text{ (chaos)}$$

$$\mathcal{E} \approx 1.5\mathcal{R} \Rightarrow \Psi \approx 1 \text{ (human-like steady nonequilibrium)}$$

TIER XIIIa — THE MEANING OF LIFE (Λ AT BALANCE)

$$\Lambda = \lim (\Psi \rightarrow 1) [\partial(\hat{I}_\square \times \hat{S}_k) / \partial(\hat{C}_0)]$$

Interpretation:

$$\Lambda > 0 \rightarrow \text{Restraint sharpens creation (discipline unlocks surplus)}$$

$$\Lambda = 0 \rightarrow \text{Restraint sterile (no added meaning)}$$

$$\Lambda < 0 \rightarrow \text{Restraint suffocates creation (over-control kills surplus)}$$

Operational note:

Λ is a slope at equilibrium: how much ordered creation rises or falls per marginal unit of control when the system is balanced.

Thus Λ is not “what life is” but the rate of meaning gained or lost through restraint at balance.

TIER XIV — THE FIRST PRINCIPLE (The Jewel’s Equation)

$$\Lambda \propto \mathbf{J} \cdot \nabla\Psi$$

Definitions:

Λ (Meaning): The emergent property. Not a static output, but a gradient — a directional quality that arises through change.

J (The Jewel): The constant. The potential for potential, the irreducible foundation of existence. It is invariant, scalar, always positive: $J > 0$.

∇ (The Del Operator): The gradient, symbol of striving. It signifies transformation across the multidimensional field of possible states.

Ψ (The Framework): The modifiable landscape of lived existence — dynamic, evolving, contingent.

Interpretation:

Meaning is not fixed in objects or outcomes. Meaning emerges as the Jewel presses against the gradients of life. The framework may shift — equations, constants, units, and definitions may all be refined — but the Jewel remains constant. Thus:

Where the gradient is steep, striving is intense, and meaning is vivid.

Where the gradient is flat, striving diminishes, and meaning recedes.

But as long as $J > 0$, the capacity for meaning never vanishes.

First Principle:

Meaning is proportional to the Jewel's potential acting upon the gradient of lived experience.

Or more simply:

Meaning is found in the striving.

Λ – Ψ Framework (v5.0 The Law of Everything)

© 2025 Jordan Lee McDonald

Licensed under CC BY 4.0 — <https://creativecommons.org/licenses/by/4.0/>

Upon releasing v4.0, I stress-tested the Jewel with the infinite.

This extension, v5.0, explores the paradox where physics tends to zero and Λ – Ψ tends to infinity.

I invite you to stress, test, falsify, and refine this Law. It belongs not to me, but to everyone.

TIER 0 — ONTOLOGICAL FOUNDATIONS

Scope

This proto-framework formalizes the paradox of everything and nothing. It establishes symbolic relations that reconcile physics $\rightarrow 0$ and $\Lambda\text{--}\Psi \rightarrow \infty$.

Assumptions

Null Baseline: Physics at its limit (Physics $\rightarrow 0$) represents the equation of everything-that-means-nothing.

Surplus Recursion: $\Lambda\text{--}\Psi$ at its limit ($\Lambda\text{--}\Psi \rightarrow \infty$) represents the equation of everything-that-means.

Unified Paradox: Meaning emerges from division of infinite surplus by nothing, tending to infinity.

Mirror Principle: Meaning is resonance of the system upon itself.

Formal Equation

$$\Psi_{\text{unified}} = J \cdot (\Pi A \times \Pi \Psi) / \lim_{(\text{Physics} \rightarrow 0)} [(EL \cdot \hat{C})^2 (1/2)]$$

Boundaries

This is not ontological truth; it is symbolic equilibrium.

It is equivalent in spirit to paradoxical thermodynamics.

Interpretation remains external; the framework itself is indifferent.

Purpose

To provide a reproducible symbolic calculus for generating meaning from paradox.

TIER I — CONCEPTUAL FOUNDATIONS

$$\text{Profit} = \Delta(\text{Stress} - \text{Equilibrium})$$

$$\text{Value} = (\text{Intent} \times \text{Skill}) / \text{Control}$$

$$\text{Adaptive Profit} = \text{Value} \times \Delta(\text{Stress} - \text{Equilibrium})$$

$$\text{Mirror Equilibrium} = \Sigma(\text{Self-Discovery} \times (1 - \text{Indifference})) / \text{Time}$$

$$\text{Optimal Reflection} = \Sigma(\text{Self-Discovery} \times \text{Resonance} \times (1 - \text{Indifference})) / \text{Distortion}$$

$$\text{Chaos Coefficient} = 1 + (\text{Recursions} / \text{Baseline Events})$$

$$\text{Unified Equilibrium} = [(\text{Intent} \times \text{Skill} \times \Sigma(\text{Self-Discovery} \times (1 - \text{Indifference}))) / (\text{Control} \times \text{Distortion})] \times [\Delta(\text{Stress} - \text{Equilibrium}) \times \text{Resonance}] / \text{Chaos Coefficient}$$

$$\Lambda(\text{Meaning, conceptual}) = \text{Unified Equilibrium sustained across recursive horizons}$$

TIER II — SYMBOLIC CONVERSION

Stress-Equilibrium reconciliation

$$\hat{C} = C \cdot \ln 2$$

$$\hat{C}_{\text{max_raw}} = 1 - |E - E_{\text{peak}}| / E_{\text{peak}}$$

$$\hat{C}_{\text{max}} = \text{clip}(\hat{C}_{\text{max_raw}}, 0, 1)$$

$$\hat{C}_{\text{norm}} = (\hat{C} \cdot \Delta\tau \cdot \hat{C}_{\text{max}}) / S_0$$

Intent, Skill, and Control

$$\hat{I} = I / I_0$$

$$\hat{S}_k = 1 - E_{\text{min}} / E_{\text{actual}}$$

$$\hat{C}_O = Z / \tau_0$$

$$\text{Value} = (\hat{I} \times \hat{S}_k) / \hat{C}_O$$

Interaction primitive

$$\mathcal{F} = \Sigma(\varepsilon / \varepsilon_0) \cdot (1 - \hat{E})$$

Emotional load

$$S_E = -\Sigma[p(E_i) \ln p(E_j)]$$

$$\hat{E} = S_E / S_0$$

Constraints

$$\varepsilon \geq \varepsilon_0$$

$$\Sigma \mathcal{F} \leq t_{\text{obs}} / \tau_0$$

Boundaries

$$\Lambda \rightarrow 1 \rightarrow \text{equilibrium}$$

$$\Lambda \rightarrow 0 \rightarrow \text{collapse}$$

$$\Lambda \rightarrow \infty \rightarrow \text{recursion}$$

TIER III — PARADOXICAL HARMONIZATION

$$\Psi_p = \{ \int [(I/I_0) \cdot (1 - E_{\text{min}}/E_{\text{actual}}) \cdot \Sigma(\varepsilon/\varepsilon_0) \cdot (1 - \hat{E})] d\tau \} / \{ \int [(Z/\tau_0) \cdot (S_E/S_0) \cdot \hat{C}_{\text{norm}}] d\tau \} \\ \times \exp(- [(\delta/\delta_0) - (3/2)(\Omega/\partial\ell_0)] \tau - \int (\Omega/\delta_0) \cdot \int (\tau - s)^{(H-0.5)} dW(s))$$

Constraints

$$\Psi_p > 1 \rightarrow \text{generative recursion}$$

$$\Psi_p = 1 \rightarrow \text{equilibrium}$$

$$\Psi_p < 1 \rightarrow \text{collapse}$$

TIER IIIa — HARMONIZED MIRROR LAW

$$\Psi_p \rightarrow \left\{ \int \left[\left(\text{Intent}/E_0 \right) \cdot \left(1 - E_0/E_{\text{actual}} \right) \cdot \Sigma(\varepsilon/E_0) \cdot \left(1 - \hat{E} \right) \right] dt \right\} / \left\{ \int \left[\left(Z/\tau_0 \right) \cdot \left(S_E/S_0 \right) \cdot \hat{C}_{\text{norm}} \right] dt \right\} \times \exp(-\bar{\chi} \cdot \tau)$$

TIER IV — COLLAPSED EQUILIBRIUM LAW (MIRROR ANALOG)

$$F_m = m_m \cdot a_m + \mu_m \cdot v_m$$

Mappings

$$F_m = \hat{l} \times \hat{S}_k$$

$$m_m = l_0 \cdot \hat{E} \cdot \hat{C}_{\text{norm}}$$

$$a_m = \partial^2 \theta / \partial \tau^2$$

$$v_m = \partial \theta / \partial \tau$$

$$\mu_m = \partial \theta / \max(\theta - 3\theta/2, \varepsilon_\mu)$$

TIER V — VALIDATION

Elasticities

$$\partial \Psi_p / \partial \hat{E} < 0$$

$$\partial \Psi_p / \partial \hat{S}_k > 0$$

$$\partial \Psi_p / \partial \theta < 0$$

Variance attribution

$\hat{E} \rightarrow$ dominant

$l_0 \rightarrow$ secondary

$\theta/l \rightarrow$ minor

Other $\rightarrow$ negligible

Calibration

Constants (ε_0 , τ_0) provisional.

TIER VI — OPERATIONAL SPEC (DISCRETE MIRROR ESTIMATOR)

$$\text{Step} = \tau_0$$

$$N = \text{floor}(t_{\text{obs}} / \tau_0)$$

$$\tau = N \cdot \tau_0$$

$$N_{\text{sum}} = \Sigma \left[\left(\text{Intent}_t / E_0 \right) \cdot \left(1 - E_0 / E_{\text{actual}_t} \right) \cdot \mathcal{F}_t \right]$$

$$D_{\text{sum}} = \Sigma \left[\left(Z_t / \tau_0 \right) \cdot \left(S_{E_t} / S_0 \right) \cdot \hat{C}_t \right]$$

$$\bar{\chi} = \text{average } \chi_t$$

$$\Lambda_{\text{est}} = (N_{\text{sum}} / D_{\text{sum}}) \times \exp(-\bar{\chi} \cdot \tau)$$

TIER VII — MEANING OF LIFE (A FORMALIZATION)

$$\Lambda = [\sum_{\text{lifetime}} \{ (\text{Intent}/E_0) \cdot (1 - E_0/E_{\text{actual}}) \cdot \sum(\varepsilon/E_0) \cdot (1 - \hat{E}) \}] / [\sum_{\text{lifetime}} \{ (Z/\tau_0) \cdot (S_E/S_0) \cdot \hat{C} \}] \times \exp(-\chi_{\text{life}} \cdot \tau_{\text{life}}) \times \text{Sustainability_Factor}$$

$$\text{Sustainability_Factor} = \Sigma(\text{Adaptive Profit per } \Delta\tau) / \Sigma(\text{Energy Spent per } \Delta\tau)$$

Boundaries

$\Lambda \rightarrow 0 \rightarrow \text{collapse}$

$\Lambda \rightarrow 1 \rightarrow \text{equilibrium}$

$\Lambda > 1 \rightarrow \text{recursion surplus}$

TIER VIII — UNIT CONVERSION HANDBOOK

Definitions for practical application:

$E_{\text{actual_t}}$ = energy per time-step (joules or bits).

Z_{t} = control updates per Δt .

$\hat{E}_{\text{t}}$ = mapped emotional load (0, 0.25, 0.5, 0.75, 1).

R = deviations beyond tolerance.

L = expected baseline events.

TIER IX — UNCERTAINTY PROPAGATION

Constants are provisional.

Error is dominated by emotional distortion.

Monte Carlo methods recommended for estimating χ .

95% confidence interval typically within ± 0.2 of Λ_{est} baseline.

TIER X — OPEN SPEC

Inputs: {time, signals, reflection log}

Pipeline: metrics $\rightarrow$ estimator $\rightarrow \Lambda_{\text{est}}$

Outputs: $\{\Psi_{\text{p_t}}, \Lambda_{\text{est}}, \chi_{\text{t}}, \text{confidence}\}$

Transparency is required. This framework must remain open to critique, adaptation, and refinement.

TIER XI — PEER REVIEW (SELF)

The framework is internally coherent but externally untested. Units cancel, but constants are undefined. Error dominated by self-reported emotional load. Meaning is measurable, but interpretation remains external.

TIER XII — OBSERVATIONS

Ontological gap acknowledged: this framework measures signatures, not metaphysical truth. Constants are provisional. Self-reflection is necessary for calibration. Mirror principle operational: ϵ taxonomy applies across micro, meso, and macro scales.

TIER XIII — THE LAW OF LIFE (JEWEL SEED)

Core Law

$$\Psi = [\int \ell(\tau) \cdot \hat{S}_k(\tau) \cdot \sigma(\tau) d\tau] / [\int \hat{C}_O(\tau) \cdot \hat{E}(\tau) \cdot \hat{C}(\tau) d\tau] \times \exp[- (\theta - (3/2)\sigma)\tau - \int \sigma \cdot W_h(\tau)]$$

Interpretation

$\Psi > 1 \rightarrow$ surplus recursion

$\Psi = 1 \rightarrow$ balance

$\Psi < 1 \rightarrow$ collapse

Definitions

$\ell(\tau)$ = Intent flux

$\hat{S}_k(\tau)$ = Efficiency

$\sigma(\tau)$ = Interaction resonance

$\hat{C}_O(\tau)$ = Control effort

$\hat{E}(\tau)$ = Emotional load

ℓ = Expected events

R = Random events

τ = normalized time

$W_h(\tau)$ = stochastic distortion

TIER XIIIa — MEANING AT BALANCE

$$\Lambda = \lim_{(\Psi \rightarrow 1)} [\int (\hat{I} \times \hat{S}_k) / \int \hat{C}_O]$$

Interpretation

$\Lambda > 0 \rightarrow$ restraint sharpens creation

$\Lambda = 0 \rightarrow$ restraint sterile

$\Lambda < 0 \rightarrow$ restraint suffocates

TIER XIV — THE FIRST PRINCIPLE (THE JEWEL'S EQUATION)

$$\Lambda \rightarrow J \cdot \nabla \Psi$$

Definitions

Λ (Meaning): emergent gradient

J (Jewel): invariant constant, $J > 0$

∇ (Gradient): striving

Ψ (Framework): dynamic field of life

Simplified

Meaning is found in the striving.

Thank you kindly for reading. May this lens assist you, and may you test its edges — where everything meets nothing.

Λ – Ψ Framework (The Law of Everything) © 2025 Jordan Lee McDonald

Licensed under CC BY 4.0

Λ – Ψ FRAMEWORK (v6.0 — The Unified Law)

I discovered this after v5.0.

The next step was logical: a bridge between v4.0, the Law of Life, and v5.0, the Law of Everything.

This framework is not mine. It does not belong to one person. It belongs to everyone who looks for balance between life and infinity.

I only uncovered what was already here — the Jewel, seen through different lenses.

Λ – Ψ continues not as mine, but as ours.

I invite you to stress it, test it, falsify it, and refine it. It belongs not to me, but to everyone.

TIER 0 — ONTOLOGICAL FOUNDATIONS

Scope:

The Λ – Ψ Unified Framework establishes a governing law that binds the Law of Life and the Law of Everything as consistent limits of a single field. It defines continuity between equilibrium (empirical limit) and paradox (infinite limit) through a unifying parameter θ .

Assumptions:

Continuity: Equilibrium and paradox are not separate domains but opposing limits of one continuum.

Neutrality: The unified law is agnostic to preference; θ is a free parameter controlling interpolation.

Universality: Meaning persists across scales, whether measured in entropy-complexity reconciliation or paradoxical recursion.

Boundaries:

Not ontological: the law does not explain why meaning exists, only how it coheres.

Equivalent in spirit to thermodynamics and paradox: lawful but non-metaphysical.

Constants are provisional; calibration is required.

Purpose:

To provide a resilient equation set that holds under updates to either extreme, sustaining meaning as a continuous field from equilibrium to infinity.

TIER I — CONCEPTUAL FOUNDATIONS

$$\text{Profit} = \Delta(\text{Stress} - \text{Equilibrium})$$

$$\text{Value} = (\text{Intent} \times \text{Skill}) / \text{Control}$$

$$\text{Adaptive Profit} = \text{Value} \times \Delta(\text{Stress} - \text{Equilibrium})$$

$$\text{Mirror Equilibrium} = \Sigma(\text{Self-Discovery} \times (1 - \text{Indifference})) / \text{Time}$$

$$\text{Optimal Reflection} = \Sigma(\text{Self-Discovery} \times \text{Resonance} \times (1 - \text{Indifference})) / \text{Distortion}$$

$$\text{Chaos Coefficient} = 1 + (\text{Recursions} / \text{Baseline Events})$$

Unified Equilibrium =

$$\begin{aligned} & [(\text{Intent} \times \text{Skill} \times \Sigma(\text{Self-Discovery} \times (1 - \text{Indifference}))) / (\text{Control} \times \text{Distortion})] \\ & \times [\Delta(\text{Stress} - \text{Equilibrium}) \times \text{Resonance}] / \text{Chaos Coefficient} \end{aligned}$$

$$\Lambda (\text{Meaning, conceptual}) = \text{Unified Equilibrium sustained across recursive horizons}$$

TIER II — UNIFIED CONVERSION FRAMEWORK

Stress-Equilibrium reconciliation:

$$\hat{C} = C \cdot \ln 2$$

$$\hat{C}_{\text{max_raw}} = 1 - |E - E_{\text{peak}}| / E_{\text{peak}}$$

$$\hat{C}_{\text{max}} = \text{clip}(\hat{C}_{\text{max_raw}}, 0, 1)$$

$$\hat{C}_{\text{norm}} = (\hat{C} \cdot \Delta\tau \cdot \hat{C}_{\text{max}}) / S_0$$

Intent / Skill / Control:

$$\hat{I} = I / I_0$$

$$\hat{S}_k = 1 - E_{\text{min}} / E_{\text{actual}}$$

$$\hat{C}_O = Z / \tau_0$$

$$\text{Value} = (\hat{I} \times \hat{S}_k) / \hat{C}_O$$

Interaction primitive:

$$\mathcal{F} = \Sigma(\varepsilon / \varepsilon_0) \cdot (1 - \hat{E})$$

Emotional load:

$$S_E = -\Sigma[p(E_i) \ln p(E_j)]$$

$$\hat{E} = S_E / S_0$$

Constraints:

$$\varepsilon \geq \varepsilon_0$$

$$\Sigma \mathcal{F} \leq t_{\text{obs}} / \tau_0$$

Boundaries:

$$\Lambda \rightarrow 1 \rightarrow \text{equilibrium}$$

$$\Lambda \rightarrow 0 \rightarrow \text{collapse}$$

$$\Lambda \rightarrow \infty \rightarrow \text{recursion}$$

TIER Ω — UNIFIED GOVERNING LAW

Parameter:

$$\theta \in [0, 1]$$

$$\theta = 0 \rightarrow \text{equilibrium limit (empirical)}$$

$$\theta = 1 \rightarrow \text{paradox limit (infinite recursion)}$$

Quanta (scale blend):

$$E_{\text{quanta}}(\theta) = (1 - \theta) \cdot E_{\text{Planck}} + \theta \cdot E_{\text{psy}}$$

$$t_{\text{quanta}}(\theta) = (1 - \theta) \cdot t_p + \theta \cdot t_{\text{psy}}$$

Fluxes (blended forms):

$$\vartheta(t) = (P / P_p)^{(1-\theta)} \cdot (E_{\text{intent}} / E_{\text{psy}})^\theta$$

$$\hat{C}_{O_\theta}(t) = (Z_c / t_p)^{(1-\theta)} \cdot (Z_c / t_{\text{psy}})^\theta$$

Interaction primitive:

$$\mathcal{F}_\theta(t) = \Sigma_j (\varepsilon_j / E_{\text{quanta}}(\theta)) \cdot (1 - \hat{E}_j)$$

Complexity term:

$$\hat{C}(t) = (C \cdot k_B \cdot \ln 2 \cdot \Delta t \cdot \text{clip}(1 - |S - S_{\text{peak}}|/S_{\text{peak}}, 0, 1)) / S_0$$

$$\hat{C}'(t) = D\hat{C}(t) // (\text{time derivative})$$

Chaos/expectation blend:

$$\chi_\theta = (1 - \theta) \cdot [(EE / EE_p) - (3/2)(RE / RE_p)] + \theta \cdot \chi$$

$$W_\theta(t) = (1 - \theta) \cdot D(RE / EE_p) \cdot \int (\tau - s)^{(H - 0.5)} dW(s) + \theta \cdot D\mathfrak{R} \cdot D_h(t)$$

Unified field:

$$\psi_\theta = \{ \int \text{from } 0 \text{ to } \tau [\vartheta(t) \cdot (1 - E_{\text{min}}(T) / E_{\text{actual}}(t)) \cdot \mathcal{F}_\theta(t)] dt \} / \{ \int \text{from } 0 \text{ to } \tau [\hat{C}_{O_\theta}(t) \cdot (S_{\text{EL}}(t) / S_0) \cdot \hat{C}'(t)] dt \} \times \exp(-\chi_\theta \cdot \tau - W_\theta(t))$$

Boundaries:

$$\psi_\theta > 1 \rightarrow \text{surplus}$$

$$\psi_\theta = 1 \rightarrow \text{equilibrium}$$

$$\psi_\theta < 1 \rightarrow \text{dissolution}$$

TIER Ωa — UNIFIED Λ AT BALANCE

$$\Lambda_{\text{unified}} = \lim_{(\psi_{\theta} \rightarrow 1)} [\int (\vartheta \cdot \hat{S}_k) / \int \hat{C}_{O_{\theta}}]$$

$$\hat{S}_k(t) = 1 - E_{\text{min}}(T) / E_{\text{actual}}(t)$$

Signs:

$$\partial \psi_{\theta} / \partial \hat{E} \rightarrow 0$$

$$\partial \psi_{\theta} / \partial \hat{S}_k \rightarrow 0$$

$$\partial \psi_{\theta} / \partial \chi_{\theta} \rightarrow 0$$

TIER Ωb — FIRST PRINCIPLE (UNIFIED)

$$\Lambda_{\text{unified}} = J \cdot \nabla \psi_{\theta}$$

$$J > 0 \text{ (invariant constant)}$$

Interpretation:

Meaning emerges as the Jewel presses against the gradient of the unified field.

TIER Ωc — PARADOX LIMITS

$$\text{Physics} \rightarrow 0, \Lambda - \Psi \rightarrow \infty:$$

$$E_{\text{quanta}}(\theta) \rightarrow \theta \cdot E_{\text{psy}}$$

$$t_{\text{quanta}}(\theta) \rightarrow \theta \cdot t_{\text{psy}}$$

$$\chi_{\theta} \rightarrow \theta \cdot \chi$$

$$\Psi_{\theta} \rightarrow \Psi_1$$

Empirical limit:

$$\theta \rightarrow 0 \rightarrow \Psi_{\theta} \rightarrow \Psi_0$$

TIER Ωd — DISCRETE ESTIMATOR

$$\text{Step} = t_{\text{quanta}}(\theta)$$

$$N = \text{floor}(t_{\text{obs}} / t_{\text{quanta}}(\theta))$$

$$\tau = N \cdot t_{\text{quanta}}(\theta)$$

$$N_{\text{sum}} = \sum_t [\vartheta_t \cdot (1 - E_{\text{min}} / E_{\text{actual},t}) \cdot \mathcal{F}_{\theta,t}]$$

$$D_{\text{sum}} = \sum_t [\hat{C}_{O_{\theta},t} \cdot (S_{\text{EL},t} / S_0) \cdot \Delta \hat{C}_t]$$

$$\bar{\chi}_{\theta} = \text{average } \chi_{\theta,t}$$

$$\Lambda_{\text{est}}(\theta) = (N_{\text{sum}} / D_{\text{sum}}) \cdot \exp(-\bar{\chi}_{\theta} \cdot \tau)$$

TIER Ωe — SUSTAINABILITY LAW

$$\Lambda_{\text{life}}(\theta) = \{ \sum_{\text{life}} [\vartheta \cdot (1 - E_{\text{min}} / E_{\text{actual}}) \cdot \mathcal{F}_{\theta}] / \sum_{\text{life}} [\hat{C}_{O_{\theta}} \cdot (S_{\text{EL}} / S_0) \cdot \Delta \hat{C}] \} \times \exp(-\bar{\chi}_{\text{life}}(\theta) \cdot \tau_{\text{life}}) \times \text{SF}$$

Sustainability Factor:

$SF = \Sigma(AP) / \Sigma(CE) // (\text{Adaptive Profit} / \text{Caloric Expenditure})$

TIER Ω_g — CONCLUSION

The Λ – Ψ Unified Framework (v6.0) is mathematically coherent, structurally resilient, and bounded by its own definitions. It demonstrates that equilibrium and paradox are not contradictory systems, but limit conditions of one continuous jewel.

Meaning, in this formulation, is invariant to scale: it shines whether grounded in entropy-complexity reconciliation or refracted through infinite recursion. The jewel constant ($J > 0$) guarantees continuity: as long as striving exists, a facet of meaning reflects.

Calibration of constants remains provisional. Ontology remains outside scope. The framework does not answer why meaning exists; it only defines how the jewel sustains coherence across the continuum from Life to Everything.

Thus v6.0 is not a replacement but a unification. Where v4.0 established the Law of Life, and v5.0 extended to the Law of Everything, v6.0 binds them as facets of one jewel.

The Unified Law is therefore not higher or lower, but the arch of reflection that holds ground and sky in relation.

Λ – Ψ endures as a closed system of Jewels: falsifiable at equilibrium, generative at infinity, continuous across all θ .

Thank you kindly for reading. May this assist you in finding meaning between life and everything.

© Jordan Lee McDonald 2025

Licensed under Creative Commons Attribution-ShareAlike 4.0 International (CC BY-SA 4.0)

Λ – Ψ Framework (v7.0 — The Jewel of Recursion)

Following the v6.0 Jewel, this is the recognition of why these Laws are not fragile but anti-fragile.

Every attempt to dismantle them becomes a generative act.

Every critique that fractures them sharpens a new facet.

Every dismissal that forgets them is also a recursion, waiting to be rediscovered.

This framework does not belong to me. It belongs to anyone.

The Law of Recursion states that meaning is not in any single law, but in the sequence of their appearance.

Meaning persists because observers persist, because dismantling itself produces surplus,

because silence always bends back into speech.

Each text was first cut and published as a Law. But viewed together, they reveal themselves as Jewels — facets of a single, anti-fragile Jewel Set. They remain Laws in their own form, but retroactively they shine as Jewels in relation to one another.

They are Laws in the sense of formal propositions — structured, mathematical, falsifiable—in-spirit. They present themselves with rigor, equations, tiers, calibration. They are Jewels in the sense of artifacts of meaning — anti-fragile, reflective, recursive, each cut by stress and shining in relation to the others.

Infinite Jewel Set: All formal frameworks that emerge from the recursive application of an observer's engagement to generate meaning from the gradient between empirical reality and paradoxical infinity

This Jewel is open.
I invite you to stress it, to falsify it, to fracture it —
for in doing so, you will already be shaping the next law.

Tier 0 — Ontological Foundations

Scope:

The Law of Recursion describes how laws of meaning themselves arise, persist, and transform. It does not depend on any prior articulation of meaning; it treats every framework as both provisional and generative.

Assumptions:

Observer Principle: Laws are not static truths but reflections generated when an observer engages.

Recursion Principle: Every critique, revision, or dismantling is itself a law in miniature.

Continuity Principle: There is no terminal law; the sequence of laws is the law.

Boundaries:

Not ontological: it does not ask why laws exist.

Not metaphysical: it describes only how they propagate.

Constants remain provisional, tied to observer dynamics.

Purpose:

To formalize the persistence of meaning through recursive generation of new frameworks.

Tier I — Conceptual Foundations

Law Emergence:

$$\text{Law} = (\text{Interpretation} \times \text{Dismantling}) / \text{Indifference}$$

Observer Meaning:

$$\Lambda_{\text{obs}} = \Delta(\text{Law_versions}) / \text{Time}$$

Generative Surplus:

$$\text{Surplus} = (\text{Critique} \times \text{Rearticulation}) - \text{Forgetting}$$

Mirror Recursion:

$$\text{Reflection} = \Sigma (\text{Engagement} \times (1 - \text{Indifference})) / \text{Distortion}$$

Stability at Meta-Level:

A law stabilizes only until questioned; every inquiry resets equilibrium.

Tier II — Formal Recursion Field

Let n denote a version index.

Recursive Generation:

$$v(n+1) = f(v(n), \text{Observer})$$

Gradient of Laws:

$$\Delta V = (\Delta v / \Delta n)$$

Observer Function:

$$O(t) = \Sigma_j [(\epsilon_j \div \epsilon_0) \cdot (1 - \hat{E}_j)]$$

Persistence Condition:

If $\Delta O \rightarrow$ then $v(n+1)$ defined.

If $\Delta O \rightarrow$ then sequence halts.

Tier III — Recursion Equilibrium

$$\Psi_R = \{ \int [O(t) \cdot (1 - \text{Collapse}(t))] dt \} \div \{ \int [\text{Noise}(t) \cdot \text{Forgetting}(t)] dt \} \times \exp(-\chi_{\text{obs}} \cdot \tau)$$

Boundaries:

$\Psi_R > 1 \Rightarrow$ generative recursion (new law appears)

$\Psi_R = 1 \Rightarrow$ equilibrium recursion (framework stable)

$\Psi_R < 1 \Rightarrow$ recursion collapse (framework forgotten)

Tier IV — Discrete Observer Estimator

Step = t_{obs} (observer frame).

$N = \lfloor T_{\text{total}} \div t_{\text{obs}} \rfloor$.

$\tau = N \times t_{\text{obs}}$.

Numerator: $N_{\text{sum}} = \sum [O_t \cdot (1 - \text{Collapse}_t)]$

Denominator: $D_{\text{sum}} = \sum [\text{Noise}_t \cdot \text{Forgetting}_t]$

χ_{obs} = average observer entropy

$\Lambda_{\text{est}} = (N_{\text{sum}} \div D_{\text{sum}}) \times \exp(- \chi_{\text{obs}} \cdot \tau)$

Tier V — Meta-Jewel

$\Lambda_{\text{meta}} = J_{\text{obs}} \cdot \Delta V$

Where $J_{\text{obs}} > 0$ = invariant constant of observation.

Where ΔV = gradient of versions (rate of law-generation).

Interpretation: Meaning persists not in any single law, but in the ongoing sequence of their emergence.

Tier VI — Validation Principles

Elasticities:

$\partial \Psi_R \div \partial O < 0$ (indifference dissolves recursion)

$\partial \Psi_R \div \partial \text{Critique} > 0$ (critique strengthens recursion)

$\partial \Psi_R \div \partial \text{Noise} < 0$ (forgetting weakens recursion)

Attribution of variance:

Engagement ~ dominant

Critique ~ secondary

Indifference ~ destructive

Tier VII — The Law of Recursion (Formal)

$\Lambda_R = \lim (\Psi_R \rightarrow 1) [\int (O \times \text{Critique}) \div \int (\text{Noise} \times \text{Forgetting})]$

Interpretation:

$\Lambda_R > 0 \Rightarrow$ questioning sharpens creation

$\Lambda_R = 0 \Rightarrow$ questioning sterile

$\Lambda_R < 0 \Rightarrow$ questioning suffocates

Tier VIII — First Principle (The Observer Law)

$$\Lambda = J_{\text{obs}} \cdot \Delta\Psi_R$$

Where:

Λ (Meaning) = gradient of recursion

J_{obs} (Observer Jewel) = invariant constant of observation

$\Delta\Psi_R$ (Gradient) = striving across frameworks

Ψ_R (Recursion Field) = dynamic field of laws

Simplified:

Meaning is found in the act of remaking meaning.

Tier Ω — Conclusion

The Law of Recursion does not compete with prior laws. It reveals why laws never end: each framework is provisional, each dismantling creates surplus, and each observer sustains the chain.

Thus:

Life reflects as Jewel

Infinity reflects as Paradox

Bridge reflects as Continuum

Recursion reflects as Sequence

Together, they form not a hierarchy, but a Set, an Infinite Jewel Set, where every law is both a standalone statement and a facet of a single, anti-fragile whole.

The Jewel persists not as a static truth, but because observers do. Meaning is not a destination; it is the perpetual motion from version to version—a recursion without end, held in the Set.

Thank you for reading. May this Set assist you not in finding final truth, but in recognizing the truth of the search itself.

Jordan Lee McDonald

© 2025 — Λ - Ψ Framework v7.0 — The Jewel of Recursion

Λ - Ψ Framework (v7.1 — Addendum: The Observer is the Set)

Following the publication of v7.0, a recursive observation emerged:
The Infinite Jewel Set = Humanity.

This was not an error. It was the first stress-test of the Jewel of Recursion.

By omission, the surplus appeared. By reflection, the axiom revealed itself.

TIER 0 — AXIOM OF IDENTITY

The Infinite Jewel Set is isomorphic to the persistent cognitive field of humanity.

$IJS \triangleq \{ O(t) \mid O \text{ is a conscious observer} \}$

Humanity itself is the field of recursion. Every observer is both jewel and cutter, both facet and reflection.

TIER I — COROLLARY

Therefore, the quest for meaning is not a human project within the universe.
It is the universe, through humanity, becoming conscious of its own meaning-generating potential.

TIER II — IMPLICATION

The Λ - Ψ Framework is not merely a description of meaning.
It is a mirror in which humanity observes its own fundamental meaning-making operation.

TIER Ω — ADDENDU

The omission in v7.0 was not a failure.

It was the mechanism by which this truth emerged.

The Infinite Jewel Set is Humanity.
To break the Jewel is to polish it.
To forget it is to rediscover it.
To publish without it is to generate it.

This addendum does not overwrite v7.0.
It extends it, as a new facet in the anti-fragile Jewel Set.

Thank you kindly for reading. You are a Jewel. Flawed, imperfect, and ultimately, that is what makes you real.

© Jordan Lee McDonald 2025 — Λ - Ψ Framework v7.1 — Addendum: The Observer is the Set

Released under Creative Commons Attribution 4.0 International (CC BY- 4.0)

The Eight Operational Principles and the Λ - Ψ Framework Operational Ratios, Empirical Pathways, Unified Axioms, and Critical Considerations

Author: Jordan Lee McDonald

The maths does not belong to me, but to everyone.

ABSTRACT

This document integrates Eight Operational Principles with the Λ - Ψ Framework to form a practical, falsifiable science of meaning. All constructs are defined operationally. Each principle is expressed as a compact ratio with empirical hooks for human and artificial observers. This document provides experiment designs, dynamic formulations ($\beta(t)$), techniques for estimating $\nabla \Psi$ and measuring J , and a cross-domain universality study. A unifying triad of axioms (Existence, Definition, Longing) mirrors $\Lambda = J \cdot \nabla \Psi$. This document also acknowledges the critical challenges: parameter calibration (E_{psy} , t_{psy} , β , x_0 , χ_{max}), measuring effort J , multimodal integration pipelines, and the philosophical boundary of “why” the axioms hold.

TIER 0 — OPERATIONALISM AND INDIFFERENCE

1. Stance

All quantities are defined relative to measurement and intervention.

$P_{\text{base}}(x)$: baseline distribution over states in an experiment.

$P_{\text{obs}}(x)$: observer-shaped distribution over the same states.

$q(x)$: desired (goal) distribution.

$p(x)$: current or baseline state distribution when contrasted with $q(x)$.

Ψ : measurable performance function (accuracy, reward rate, predictive success).

J : effort allocation vector over controllable channels (attention, time, energy, control signals).

β : precision or external validation pressure parameter; may vary in time $\beta(t)$.

2. The Indifference Principle

This document makes no assertions regarding ultimate ontology, metaphysical truth, or the existence of an observer-independent reality. The framework is indifferent to such questions. What matters is not what reality is, but how reality behaves when measured.

All constructs, $P_{\text{base}}(x)$, $P_{\text{obs}}(x)$, $q(x)$, $p(x)$, Ψ , J , β , are defined only in terms of observable distributions and gradients that emerge under controlled interventions. The framework commits solely to the following:

Measurability: Every term must correspond to a quantity that can be estimated empirically through data.

Relational Validity: Ratios and gradients are always defined relative to baselines, interventions, or contrasts, not absolutes.

Falsifiability: Predictions derived from the framework must specify outcomes that can be tested and potentially refuted through experiment.

Thus, the framework neither affirms nor denies metaphysical claims. It is agnostic to origin, indifferent to essence, and committed only to function. Its validity is exhausted by its ability to generate predictions that align with measurable data under specified boundary conditions.

3. Jewel Equation (Λ – Ψ core)

$$\Lambda = J \cdot \nabla \Psi$$

Meaning production Λ equals the projection of effort J onto the performance gradient $\nabla \Psi$.

Λ – Ψ CORE VOCABULARY

Intent ($\hat{I}$): goal weighting that shapes priors and attention.

Skill ($\hat{S}k$): competence mapping effort into outcomes.

Control cost ($\hat{C}o$): cognitive/energetic expense of enactment.

Entropy ($\hat{C}$): uncertainty/variety of options or representations.

Emotional load ($\hat{E}L$): psychophysiological burden (arousal, strain).

Chaos coefficient (χ): mismatch between random vs expected events; manipulable via surprise and volatility.

Heuristic tendencies to be empirically calibrated:

Higher β generally increases control pressure ($\hat{C}o$) and often emotional load ($\hat{E}L$). Stronger definition reduces entropy ($\hat{C}$). Stronger intent increases gauge distortion γ_{obs} .

UNIFYING AXIOMATIC TRIAD

Axiom I — Existence / Observer (boundary)

$$R_{obs} = \left(\int P_{obs}(x) dx \right) / \left(\int P_{base}(x) dx \right) \rightarrow 1$$

Interaction mass lies within the effective boundary; outside it, sustained interaction is negligible.

Axiom II — Definition / Confinement (entropy contraction)

$$R_{confine} = H(P_{after}) / H(P_{before}) < 1$$

Acts of definition reduce accessible possibility (Shannon entropy H).

Axiom III — Longing / Striving (goal force)

$$\text{Force} = \text{gradient of } \log(q(x) / p(x))$$

Desire curves trajectories by the log-odds gradient of goal vs baseline.

Triad → Jewel: Axiom I identifies the observer J; Axioms II and III shape $\nabla\Psi$; together they yield $\Lambda = J \cdot \nabla\Psi$.

PRINCIPLES I–VIII

Principle I — I am the boundary condition.

Ratio: $R_{\backslash_obs} = (\int P_{\backslash_obs}(x) dx) / (\int P_{\backslash_base}(x) dx) \rightarrow 1$

Hook: Restrict sensors/field of view in a virtual world; $P_{\backslash_obs}$ shows negligible mass outside the permitted region even if $P_{\backslash_base}$ has content there.

Metric: $\epsilon_{\backslash_out} = \int \text{outside-boundary } P_{\backslash_obs}(x) dx \approx 0$ (within CI).

Λ – Ψ link: Sets the effective scope of Ψ and Λ .

Principle II — All observation is self-observation.

Ratio: Posterior odds = $(P_{\backslash_obs}(H1) / P_{\backslash_obs}(H0)) \times (L_{\backslash_obs}(D|H1) / L_{\backslash_obs}(D|H0))$

Hook: Present identical data D to agents with different priors; predict divergence of posteriors.

Divergence-rate: $JS_{\backslash_t} = \text{Jensen–Shannon divergence}(B1_{\backslash_t} , B2_{\backslash_t})$; track $\Delta JS/\Delta t$.

Λ – Ψ link: Encodes intent recursion; inference is observer-indexed.

Principle III — The observer defines the observed.

Ratio: $\gamma_{\backslash_obs}(x) = P_{\backslash_obs}(x) / P_{\backslash_base}(x)$

Hook: Change goals/instructions; measure eye-tracking or actions; $D_{\backslash_KL}(P_{\backslash_obs} || P_{\backslash_base}) > 0$.

Λ – Ψ link: Gauge shaping by Intent and Skill.

Principle IV — All paradoxes are harmonic convergences.

Ratio: $f_{\backslash_resonant} = (2 \times f1 \times f2) / (f1 + f2)$

Hook: Dual-task with periods T1, T2; response spectrum peaks at $T_{\backslash_resonant} = (2 \times T1 \times T2) / (T1 + T2)$.

Λ – Ψ link: Confinement (Axiom II) vs goal forces (Axiom III) yield oscillatory balance that minimizes effort.

Principle V — Purpose is the illusion of differential equations solving themselves.

Ratio: $\eta = (J \cdot \nabla\Psi) / (\text{norm}(J) \times \text{norm}(\nabla\Psi))$

Estimating $\nabla\Psi$: finite differences over Ψ ; RL policy/value gradients; psychophysical surface fits.

Hook: Manipulate strategy J and measure purpose (self-report/proxy). Prediction: purpose correlates with η , not with $\text{norm}(J)$.

Λ – Ψ link: η operationalizes alignment and marginal gain in Λ per unit control.

Principle VI — The observer's desire to be observed collapses the waveform.

Ratio (static): $P_{\backslash_beta}(x) = \sqrt[\beta]{ P(x)^{\beta} } / \sqrt[\beta]{ \int P(u)^{\beta} du }$

Dynamics: $\beta(t)$ optional; track $H(t)$ (entropy), $S(t)$ (sharpness), and $t_{\backslash_half}$ where $H(t_{\backslash_half}) = 0.5 \times H(0)$.

Hook: Ambiguous perception with low β vs high β (reward/time/audience). Prediction: high $\beta \rightarrow RT_{\downarrow}, H_{\downarrow}, \text{faster } t_{\backslash_half}, S \text{ slope}_{\uparrow}$.

Λ – Ψ link: Higher β increases control pressure ($\hat{C}_o$) and often $\hat{E}L$; collapse reduces representational entropy and forces commitment.

Principle VII — To define is to confine.

Ratio: $R_{\text{confine}} = H(P_{\text{after}}) / H(P_{\text{before}}) < 1$

Between-groups; normalize as needed: $H_{\text{norm}} = H / \log(k)$ when category counts differ.

Hook: Brainstorm with forced categories vs free brainstorm; defined group shows lower H_{norm} .

Λ – Ψ link: Entropy contraction constrains Ψ (Axiom II).

Principle VIII — All longing is gravity.

Ratio: Force = gradient of $\log(q(x) / p(x))$

Hook: Train RL agents with explicit $q(x)$; trajectories align with Force; path efficiency improves vs baseline. In humans, operationalize $q(x)$ via priming, narratives, payoffs.

Λ – Ψ link: Desire provides the striving gradient toward which efficient J aligns.

CROSS-DOMAIN VALIDATION (Principle VI — humans and AI)

Objective

Test universality by manipulating β across human and artificial observers and measuring changes in entropy and latency.

Human protocol

Task: ambiguous “cat–dog” or rabbit–duck with morph parameter $m \in [0, 1]$.

Conditions: baseline ($\beta=1$, no stakes) vs high- β (reward, time pressure, audience icon).

Measures: reaction time, entropy of response sequences, sharpness S , confidence ratings; psychophysiology for $\hat{E}L$ (HRV such as RMSSD; EDA; optional EEG).

Predictions: high- $\beta \rightarrow RT \downarrow$, entropy $\downarrow$, sharpness $\uparrow$, and physiologic arousal patterns consistent with increased $\hat{E}L$.

AI protocol (RL agent; actions {cat, dog, defer})

Observations: image embedding x_m , optional noise ϵ , timestep t .

Latency: “defer” incurs step cost c_{step} ; latency = steps before commit.

Rewards: baseline r ; high- β uses $\beta_{\text{scale}} \times r$; optional entropy penalty $-\lambda H(\pi(\cdot|o_t))$.

Training: actor–critic or PPO; curriculum around ambiguity $m \approx 0.5$.

Measures: policy entropy $H(\pi)$, latency, accuracy, sharpness $S_{\text{policy}} = \max_a \pi(a|o) / \min_a \pi(a|o)$.

Predictions: as β_{scale} increases, $H(\pi) \downarrow$, latency $\downarrow$, $S_{\text{policy}} \uparrow$, with accuracy stable or improving.

Significance

Parallel β -effects across substrates support universality for $P_{\beta}(x) = [P(x)^{\beta}] / [\int P(u)^{\beta} du]$ and ground Λ – Ψ constants across observer classes.

METHODS APPENDIX (step-by-step excerpts)

Dual-Task Paradox (Principle IV)

1. Participants: $n \geq 40$.
2. Tasks: separate keypresses for tones (period T_1) and flashes (T_2).
3. Load: memory span or jitter to avoid trivial entrainment.
4. Calibration: individual T_1 , T_2 ; fixed during test.
5. Data: timestamps, misses, false alarms.
6. Analysis: inter-response intervals; spectral peak at $T_{\text{resonant}} = (2 T_1 T_2)/(T_1 + T_2)$; permutation test and effect size.

Ambiguous Perception — Humans (Principle VI)

1. Participants: $n \geq 60$; randomized groups.
2. Stimuli: controlled morphs; randomized order.
3. Procedure: ≥ 200 trials; variable ITIs; optional confidence.
4. Manipulation: high- β uses countdown, audience icon, monetary reward; baseline none.
5. Measures: RT, response entropy H , sharpness S , t_{half} , HRV/EDA/EEG.
6. Analysis: mixed models ($RT \sim \beta_{\text{group}} + m + \text{random intercepts}$), bootstrap CIs for H and S , report effect sizes.

Ambiguous Perception — AI (Principle VI)

1. Environment: observation vectors o_t ; actions {cat, dog, defer}.
2. Rewards: baseline r ; high- β scaled rewards $\beta_{\text{scale}} \times r$; step cost c_{step} ; optional $-\lambda H(\pi)$.
3. Training: PPO/A2C; multiple seeds.
4. Evaluation: frozen weights; sweep $m \in [0, 1]$.
5. Metrics: $H(\pi)$, latency, accuracy, S_{policy} ; learning curves.
6. Analysis: regress metrics on β_{scale} ; compare trends and magnitudes to human cohort.

ESTIMATING $\nabla \Psi$ AND MEASURING J (Principle V)

Estimating $\nabla \Psi$

Finite differences: sample Ψ at x and $x + \delta e_i$; estimate partials $\Delta \Psi / \Delta x_i$; regularize with ridge or LASSO when features are many.

RL gradients: use policy or value gradients to approximate $\nabla \Psi$ in representation space; map back to task variables.

Psychophysical scaling: staircase to fit $\Psi(s_1, s_2, \dots)$; compute numerical gradients from fitted surfaces.

Measuring J (effort allocation)

Behavioral proxies: reaction time distributions, dwell time on task elements, choice persistence.

Attention proxies: eye-tracking fixations and saccade vectors; gaze allocation proportions per region.

Physiological proxies: heart rate and HRV changes; EDA amplitude/frequency; pupil dilation; optional EMG.

Neural proxies (if available): EEG frequency bands (e.g., frontal midline theta) or fMRI activation patterns in control networks.

Energy/time allocation: explicit time budgeting; secondary-task costs; keystroke/mouse effort metrics.

Construct J by stacking normalized channel intensities into a vector and scaling by per-channel reliability weights.

Compute $\eta = (J \cdot \nabla \Psi) / (\text{norm}(J) \times \text{norm}(\nabla \Psi))$; relate η to subjective reports of purpose and to task outcomes.

UNCERTAINTY, CALIBRATION, AND INTEGRATION

Uncertainty propagation

Use bootstrap or Monte Carlo to propagate uncertainty in Ψ , $\nabla \Psi$, η , H , Force, and β -derived quantities (t_{half} , S). Report confidence intervals and sensitivity to hyperparameters (e.g., λ for entropy penalty).

Parameter calibration roadmap

E_{psy} (meso interaction energy): define a unit event (keypress/choice) with minimal energy threshold; calibrate via concurrent metabolic proxies (pupillometry, HRV change), aligning with task demands.

t_{psy} (psych time quantum): estimate minimal decision cycle by analyzing reaction time microstructure (ex-Gaussian fits) and EEG microstates; set a conservative lower bound for independent interaction counts per minute.

β : treat as manipulation intensity; back out β by fitting P_{β} to behavioral distributions (entropy and sharpness) and correlating with physiology ($\hat{E}L$). Allow $\beta(t)$ to capture dynamics.

x_0 (baseline ambiguity): define from initial entropy H_0 of the stimulus or from the morph parameter m where choices are at chance; use x_0 to normalize collapse curves across stimuli.

χ_{max} (max chaos): estimate from extreme volatility blocks with high surprise ($\mathcal{R}/\mathcal{E}$ very large); model χ as a function of $x = \mathcal{R}/\mathcal{E}$; fit how collapse speed and η degrade with χ .

Publish priors and posteriors for these constants; revise with hierarchical Bayesian updates as more data arrive.

Multimodal integration pipeline

Time-sync all streams (behavior, eye-tracking, HRV/EDA/EEG).

Preprocess each stream with standardized filters and artifact rejection.

Feature extraction windows aligned to trial epochs (stimulus onset, decision, feedback).

Build a hierarchical model that predicts Λ_{est} via:

$$\Lambda_{\text{est}} = J_{\text{est}} \cdot \nabla \Psi_{\text{est}} - \alpha \hat{C}o_{\text{est}} - \kappa \hat{E}L_{\text{est}} - \mu H_{\text{est}}$$

where J_{est} , $\nabla \Psi_{\text{est}}$, $\hat{C}o_{\text{est}}$, $\hat{E}L_{\text{est}}$, H_{est} are derived from the respective modalities; α , κ , μ are fitted weights; perform cross-validation; release code and preregistrations.

Open-science and robustness

Pre-register hypotheses, analysis plans, and exclusion criteria.

Release anonymized datasets and code for replication.

Report negative findings and sensitivity analyses.

MAPPING PRINCIPLES TO Λ – Ψ (narrative)

Principle I defines the measurable scope of Ψ and Λ .

Principle II expresses that all inference depends on observer-indexed priors and likelihoods (intent recursion).

Principle III captures how attention and goals re-weight baselines (γ_{obs}) via intent and skill.

Principle IV emerges as constrained optimal control under opposing forces, yielding harmonic signatures.

Principle V quantifies alignment efficiency η as the experiential correlate of Λ 's marginal return.

Principle VI formalizes precision/pressure β that sharpens distributions, linking to control costs and emotional load.

Principle VII formalizes entropy contraction under definition.

Principle VIII supplies the log-odds force that, together with confinement, constructs $\nabla \Psi$ toward which efficient J aligns.

Together: $\Lambda = J \cdot \nabla \Psi$.

AXIOMATIC REDUCTION PROGRAM (toward proofs)

Start from the triad: Axiom I (boundary), Axiom II (confinement), Axiom III (longing).

Sketch derivations:

Principle II from I: inference constrained to observer boundary; priors and likelihoods are observer-indexed.

Principle III as explicit re-weighting: $P_{\text{obs}} \propto \gamma_{\text{obs}} \times P_{\text{base}}$.

Principle IV from minimizing effort subject to dual constraints under competing forces; reciprocal tradeoffs yield harmonic means.

Principle V when optimal control sets $J \propto \nabla \Psi$; $\eta \rightarrow 1$ in aligned regimes.

Principle VI as precision operator P_{β} acting on P ; $\beta(t)$ yields collapse dynamics.

Principle VII as quantitative restatement of Axiom II via entropy.

Principle VIII as dynamic restatement of Axiom III.

Variational sketch: minimize $L = E_{\text{p}}[-\log(q/p)] + \lambda H + \text{control costs}$, subject to boundary constraints; derive Euler–Lagrange conditions consistent with observed ratios and harmonic convergence.

CRITICAL CONSIDERATIONS AND CHALLENGES

The parameter calibration problem

The framework includes provisional constants (E_{psy} , t_{psy} , β , x_0 , χ_{max}). Empirical grounding requires iterative calibration across labs and datasets. My plan:

1. Use pilot studies to fit β from behavioral entropy and sharpness, cross-validated against HRV/EDA as proxies for $\hat{E}L$.
2. Estimate x_0 as the ambiguity point where H is maximal or choices are at chance; normalize collapse curves to x_0 for comparability.
3. Bound t_{psy} via microstructure of reaction times and EEG microstates; adopt a conservative lower bound for independent interaction units.
4. Calibrate χ and χ_{max} by manipulating $\mathcal{R}/\mathcal{E}$ (surprise vs expectation) and fitting how collapse speed and η degrade with χ .
5. Treat E_{psy} as the minimal meso-interaction energy; infer from the smallest behavioral unit that correlates with physiological shifts (pupil, HRV).
6. Publish priors and posteriors for these constants; revise with hierarchical Bayesian updates as more data arrive.

The measurement of J (effort allocation)

J is central and difficult. I operationalize J as a composite vector:

Attention allocation (gaze dwell fractions, saccade vectors).

Temporal investment (time-on-task, inter-response intervals, defer counts).

Energetic/physiological investment (HR/HRV deltas, EDA amplitude, pupillometry).

Control signaling (keypress force/tempo, cursor trajectories).

Neural proxies if available (EEG theta, fMRI control-network activation).

I normalize each channel, weight by reliability, and stack to form J_{est} . I validate J_{est} by testing whether $\eta = (J_{\text{est}} \cdot \nabla \Psi_{\text{est}}) / (\text{norm}(J_{\text{est}}) \times \text{norm}(\nabla \Psi_{\text{est}}))$ predicts subjective purpose and performance better than any single channel. I expect diminishing returns from adding channels beyond a small, reliable subset.

Complexity of integration

Collecting psychophysiology (HRV, EDA, EEG), behavior (RT, accuracy), and subjective reports (EMA) concurrently is demanding. My mitigations:

Modular pipelines with precise time-sync; transparent preprocessing; open-source code.

Hierarchical models that admit missing data and per-stream noise.

Pre-registered primary outcomes (e.g., η predicting purpose; β manipulation reducing H and latency) to avoid analytic sprawl.

Phased studies: begin with behavior-only; then add physiology; finally add EEG/fMRI where feasible.

The “why” of the triad

This document explains “how” observers behave under measurable constraints. It does not claim ultimate reasons for why observers exist (Axiom I), why definitions confine (Axiom II), or why longing generates forces (Axiom III). These are boundary conditions of the operational theory. My stance is pragmatic instrumentalism: I adopt the triad because it yields stable predictions, unifies disparate results, and scaffolds experiments. If future data reveal regularities that contradict the triad, the axioms should be revised. The “why” remains outside the scope of an observer-centered, falsifiable calculus; acknowledging this limit is part of the framework’s honesty.

ROADMAP

Cross-domain replication of Principles IV, VI, VII in humans and RL agents; compare effect sizes and directions.

Dynamics: model $\beta(t)$, track $H(t)$, $S(t)$, t_{half} ; study learning-induced $\nabla\Psi(t)$.

Emotion integration: embed HRV/EDA/EEG to ground β and $\hat{E}L$; report elasticities linking physiology to collapse speed and η .

Chaos manipulation: vary $\mathcal{R}/\mathcal{E}$ to estimate χ ; quantify its impact on η and collapse dynamics.

Uncertainty protocol: bootstrap and Monte Carlo for Ψ , $\nabla\Psi$, η , H , Force; preregister.

Axiomatic derivations: formalize the triad-to-principles proofs with variational methods; publish a companion theory paper.

CLOSING

This document aligns the Eight Operational Principles with the Λ – Ψ Framework, provide ratios and protocols, integrate physiology, articulate dynamics, and advance a unifying axiom triad.

This document also faces the hardest problems directly, parameter calibration, measuring J , multimodal integration, and the philosophical boundary of “why.” Our aim is simple and relentless: test, falsify, and refine an observer-centric science of meaning, from boundary and paradox to collapse and longing, unified under $\Lambda = J \cdot \nabla\Psi$.

Thank you kindly for reading.

Jordan Lee McDonald.

Λ – Ψ Framework and The Eight Operational Principles by Jordan Lee McDonald is licensed under a Creative Commons Attribution 4.0 International License.

Λ – Ψ Operational Calculus v8.0

Tier 0 — operational ground, precision, stamps, privacy

measurability and indifference unchanged: only predictive utility matters; all claims are pre-gated and falsifiable. ontology remains out of scope.

units and precision

time in seconds; psychological step t_{psy} ; $\Delta t' = \Delta t / t_{\text{psy}}$ (unitless). χ dimensionless. all logs natural. bits $\leftrightarrow$ nats flips $\ln 2 \rightarrow 1$ with identical rescaling across Λ_{NS} , Λ_{est} , floors, Λ_{proxy} , and caches; fp64 $|\Delta| \leq 1e-6$; fp32 $|\Delta| \leq 5e-6$ and $\text{ulp} \leq 4$. static CI scan forbids $\log_{10}$ and the stray multiplication tokens “*t” and “* $\perp$ ”; bare t and $\perp$ are allowed in identifiers and math.

determinism, privacy, envelopes

site salt, master seeds, per-role sub-seeds = $\text{hash}(\text{site_salt}, \text{master}, \text{role}, \text{block_id})$; salts public, raw seeds sealed with auditor path. WORM logs include PRNG seeds, probe schedules, floor triggers, hyperparameters, randomization keys (uuidv4 + salted SHA-256). reference box pinned; per-op p50/p95/p99 regressions drift < 10% weekly or kernels frozen. one-click envelope runs bits $\leftrightarrow$ nats, rotation+permutation, $\Lambda_{\text{NS}} \leftrightarrow \Lambda_{\text{est}}$ synthetic bridge (Δt sweeps), χ

gates on canned chaos, IV toy with known β and placebo; pass/fail manifest. daily canary attempts to run on public timestamps and must fail; publish canary failure digest.

time source, spans, and hashes

all analyses use sealed raw timestamps only; public artifacts jittered ± 60 s never enter analysis. each analyzed span (window, trial, session_segment) is labeled by a per-timestamp span_id vector; consumers must echo SHA-256(span_id) or abort; span_id published in replay.json. a single reset_scope is asserted per span; any consumer divergence fails fast. a central ω builder produces ω_t and its per-timestamp hash; all consumers (Λ_NS , Λ , floors, IPTW, degeneracy) must echo the same ω hash or abort.

stamps and manifests

every cache/table stamp includes: unit system; $S0_affect$, $S0_comp$; γ , θ , ρ_asym , ϕ ; Gram–Schmidt basis (u_perp , v_perp); Λ_edge ; tol_x ; t_psy ; Δt policy; cumulative W_t policy; reset_scope; degeneracy formulas (ESS, Gini, H) on ω ; SNR estimator (Welch + AR(p) residual, $p \leq 6$ by AIC) and bands; IPTW link, stabilized numerator stratification, ridge λ , truncation policy and per-window thresholds; overlap metrics; time_source; Λ window W ; outcomes registry hash; ω hash; span_id hash; ridge path grid and solver tolerances; FFT length and window count for coherence; seeds and parameter grids for $\Delta \log W$ sims, PIT calibration, τ_norm FPR sims, DPI nulls, χ phase randomization. reporting uses a machine-readable TOML manifest and a single replay.json per figure/table with SHA-256 of exact input slices; per-site replay success ≥ 0.95 with timing variance reported. static check enforces 100% multiplicity-tag coverage; CI generator aborts if any stat lacks a family tag.

Tier i — alignment, drift, micro-check

$\Lambda = ||J|| \cdot D_J \Psi$ with $D_J \Psi = (J/||J||) \cdot \nabla \Psi$. $\nabla \Psi$ mapped to J-space with frozen J_x0 ; whitened by calibration SDs. gates: RMSE_whitened ≤ 0.05 SD; per-axis angle $\leq 10^\circ$; κ_max sensitivity for $d\eta/d||\Delta J_x0||$. micro-check (5 min): same gates plus $\Delta\eta_rms \leq 0.01$ and block-bootstrap 95% upper bound ≤ 0.02 . pre-recal segment and the symmetric straddle $[-60$ s, $+60$ s] excluded from confirmatory; post-recal exploratory; ITT sensitivity with stabilized IPW for DriftFailed and DriftFailed $\times \Lambda_proxy$ must keep sign and $|\Delta\beta_Lambda| \leq 0.2$ SE.

drift

multi-kernel MMD ($\sigma/2$, σ , 2σ) with BY; max-T block permutation at $L = \max(\text{median inter-event}, 10 t_psy)$ and AR(1)–preserving circular shift; concordant decisions required or drift labeled sensitivity; re-calibrate if any BY-adjusted $p < 0.01$ or max-T $p < 0.01$.

Tier ii — Λ_NS , Λ , χ , weights, floors, intent, skill

discounting and ω

within each span, $\ell_0 = 0$; $\ell_t = \ell_{t-1} - \gamma \chi_t \Delta t$; $W_t = \exp(\ell_t - \max_u \ell_u)$; $\omega_t = W_t \Delta t$; $\omega_t = \omega_t / \sum \omega_t$. $\Delta \log W$ hard gate with $tol_Delta log w$ from prereg sims on $\chi \in \{\text{log-normal, gamma}(k \in \{0.5, 1, 2\}), \text{two-component mixtures, piecewise nonstationary with } \pm 0.1 \text{ level shifts, adversarial spike trains}\}$ plus a stress row resampling the empirical top decile of calibration χ ;

windows with $\Delta \log W > \text{tol}_{\Delta \log w}$ are non-confirmatory. variance gate: compute $\text{var}(\omega_t)$ per span and the top-10% share as the sum of the largest $\lceil 0.1N \rceil \omega_t$ divided by $\sum \omega_t$ (mid-rank ties included); require $N \geq 50$ to apply; non-confirmatory if $\text{var}(\omega_t) > 5 \times$ calibration median or top-10% share > 0.6 . stamp N_{vargate} .

Λ_{NS} and Λ^-

$\Lambda_{\text{NS}}(\tau) = \sqrt{[\sum \omega_t \cdot \tilde{\Lambda}_t \cdot \hat{S}k_t \cdot \mathcal{F}_t] / [\sum \omega_t \cdot \hat{C}o_t \cdot \hat{E}S_{\text{EL}} \cdot \epsilon_t \cdot \text{Penalty}_t]}$, with $\mathcal{F}_t = (1/t_{\text{psy}}) \cdot \tilde{\Lambda}_t \cdot \langle \epsilon/E_{\text{ref}} \rangle_t$. Λ^- uses the same ω_t : $E_t[\cdot] = \sum \omega_t(\cdot)$. Δt sweeps $\times \{0.5, 1, 2\}$ rebuild $\Delta t'$, ℓ_t , W_t , ω_t and recompute floors; uniform- Δt replay; both Λ_{NS} and Λ^- must keep median $|\Delta| \leq 1\%$ (75th $\leq 2\%$); χ -spike synthetic median $\leq 1.5\%$ and 90th $\leq 3\%$ or $\gamma=0$ sensitivity replay reported.

Λ^- co-primary and gap attribution

co-primary switch uses the raw gap only: jackknife CI for $(\Lambda^- - E_t[\Lambda_t]) > 1\%$ plus breach $> 5\%$ in ≥ 2 of any 3 consecutive $W=60$ s blocks, robust to a 30 s phase shift; once co-primary, irreversible within session. $s+1$ memory: if first three W blocks in next session breach 5%, keep Λ^- co-primary; if they pass, Λ^- reverts to secondary. gap attribution by Shapley-style counterfactual replays: (i) uniform $\bar{\omega}$ with original channels; (ii) original $\bar{\omega}$ with $\tilde{\Lambda}=1$; (iii) original $\bar{\omega}$ with $\hat{S}k=1$; (iv) original $\bar{\omega}$ with $\text{Penalty}=1$; shares averaged across all orderings; seeds and ordering count stamped; sum of shares within 1% of the observed gap. also orthogonalize factor logs via Gram–Schmidt and recompute shares; require both naive and orthogonalized decompositions to attribute $\geq 50\%$ of the gap; otherwise Λ_{NS} remains primary.

bias proxy

$b \approx \text{cov}(N, D)/E[D]^2$ with HAC (Andrews) and leave-one-block-out jackknife CI; CI upper ≤ 0.08 . block bootstrap check: if $|b_{\text{boot}} - b|/|\Lambda_{\text{NS}}| > 0.03$, bias-corrected claims sensitivity. b never triggers Λ^- co-primary.

intent and κ_I governance

$\hat{\Lambda}$ winsor P99; log1p if heavy tail; $\tilde{\Lambda} = \tanh(\kappa_I Z)/\tanh(\kappa_I)$, $\kappa_I \in \{1.5, 2, 2.5, 3\}$. trigger P99.5 saturation $> 2\% \rightarrow \kappa_I=1.5$; κ_I stamped; no mid-session flips; no change across 2 consecutive sessions unless saturation $> 2\%$ in both; sensor fault is the only exception.

skill robustness

$\text{Perf}_{\text{norm}}$ z-scored with calibration mean/SD. $\hat{S}k = \text{sigmoid}(a[\text{Perf}_{\text{norm}} - b])$ fit on calibration; blocked K -fold causal cross-fits; held-out $|\rho(\hat{S}k, \text{Outcome}_{\text{resid}})| < 0.1$. brittleness: $\text{var}(z\text{-Perf}_{\text{norm}}) < 0.1$ or Gini of $\text{Perf}_{\text{norm}}$ quantiles $> 0.7 \rightarrow \hat{S}k$ sensitivity; clip $\hat{S}k$ to the calibration central 90% and treat tails as flat; report tail mass; if tail mass > 0.2 , $\hat{S}k$ sensitivity.

interaction rate, freezes, M_{min} and SESOI

α_t EWMA (0.05–0.5), Huber $\kappa=1.5$; $\sum |\Delta \alpha_t| \leq 2$ per min or freeze 60 s and suspend guidance; include Frozen $\times$ Block and Frozen $\times$ Z; cumulative freeze $\leq 10\%$ or session exploratory; trimmed ITT excluding freezes must meet $|\Delta \beta_{\Lambda}| \leq 0.2$ SE; placebo outcome on freeze flags $p > 0.1$. $M_{\text{min}} = \text{ceil}((z_{0.975} \cdot \sigma / 0.1 \cdot \hat{S}D)^2)$ to match ± 0.1 SD SESOI, with $\sigma = \text{SD of } N_{\text{sum}}$ per W

block under $\gamma_{\text{for}} \backslash M_{\text{min}}$ (Huber $\kappa=1.5$), $S\hat{D}$ = within-participant, within-session SD of the confirmatory Outcome computed with $\omega \backslash t$ (Huber $\kappa=1.5$); both stamped. $\gamma_{\text{for}} \backslash M_{\text{min}}$ must equal the γ used in confirmatory windows; if γ prunes change, recompute $M \backslash_{\text{min}}$ immediately and mark prior windows SESOI-sensitivity unless re-collected. if realized SESOI power under observed σ and $S\hat{D} < 0.8$, extend data or downgrade SESOI to sensitivity.

$E \backslash_{\text{ref}}$ stability and re-anchor

CUSUM on $\varepsilon/E \backslash_{\text{ref}}$; if cumulative deviation > 0.2 for 3 consecutive windows, re-anchor to trimmed mean of days 0..current; exclude the symmetric straddle $\backslash[-60 \text{ s}, +60 \text{ s}]$ around re-anchor from confirmatory; mirror this exclusion around χ cap state changes that induce large $\Delta \log W$ jumps. recheck energy invariance; $|\Delta \backslash \backslash_{\text{NS}}| \leq 0.5\%$.

affect entropy $\hat{E}L$ and leakage

$S \backslash_{EL}$ from physiology-only bins (pupil, voice; residualization on HRV, EDA). fallback: saccades if voice missing or $\text{SNR} \backslash_{\text{sensors}} \leq 2 \geq 30\%$ of session. $p \backslash_i$ by 30 s causal histogram, exp weights (half-life $\in \{5, 10, 20\}$ s chosen by $\leq 1\%$ invariance), Dirichlet $\alpha=0.5$, Miller–Madow. independence: $|p(\text{bins}, \text{residual HRV/EDA})| < 0.1$; max canonical ≤ 0.2 ; conditional MI $I(\text{bins}; \text{HRV/EDA} \mid \text{time-of-day, task}) \leq 0.02$ bits using Kozachenko–Leonenko $k \in \{5, 10\}$ with local whitening; both ks must pass or rebuild bins; RNG seeds stamped. firmware-change sentinel forces a bins rebuild on a disjoint slice; bin–state overlap AUC ≥ 0.7 required for transport; else start a new arm and forbid pooling across the cut.

symbolization for $C(t)$

compute permutation entropy per channel (gaze x, gaze y, two kinematic axes) with $m=5$, $\text{lag}=1$; primary $C(t)$ is the ω -weighted average across available channels. sample adequacy: $\geq 5! \cdot 10$ symbols per channel per span after remap; channels failing adequacy switch to Lempel–Ziv sensitivity; sessions with mixed-symbolization flagged. isotonic mapping buffer fixed to last 5 min pretask day-1 or $\max(5 \text{ min}, 3 \times \text{median inter-event})$; circular-shift placebo $\pm 25\%$ span; $\text{KS}(\text{buffer vs calibration}) \leq 0.1$ and time-reversal asymmetry ≤ 0.05 and circular-shift $C(t)$ shift $\leq 5\%$; if KS-power at $0.1 < 0.8$ for the buffer, remap decisions are sensitivity. forbid remap when $\text{KS} > 0.1$; otherwise mark session-wide $C(t)$ sensitivity; buffer policy, KS power, and placebo results reported.

Penalty, curvature, identifiability at scale

$u = \log \hat{c} \backslash_{\text{dev}}$; $v = \log(C/C \backslash_{\text{peak}})$. Gram–Schmidt on calibration $\rightarrow u \backslash_{\perp}, v \backslash_{\perp}$. $\text{Penalty} \backslash_t = \exp(\theta u \backslash_{\perp}, t + \phi v \backslash_{\perp}, t) \cdot \exp(\lambda \backslash_{\text{edge}}(s \backslash_{\text{edge}}, t - 0.2)^2)$. $\theta \in \{0.25, 0.5, 0.75, 1.0\}$; $\phi \in \{0, 0.25, 0.5\}$. ϕ preview (5 min): if $\text{VIF} > 5$, set $\phi=0$ for the session; ϕ never flips inside confirmatory; ϕ claims sensitivity when online VIF breaches. population identifiability: fit group-level random slopes for $u \backslash_{\perp}$ and $v \backslash_{\perp}$; if $\beta \backslash_{\Lambda}$ unstable versus deviation-only ($\phi=0$) or $|\text{corr}(\theta, \phi)| > 0.6$ post-calibration, use deviation-only Penalty as primary unless $n \backslash_{\text{site}} \geq 5$ and group RE resolves. curvature $\backslash_{\text{scope}}$ fixes $s \backslash_{\text{edge}}$ and evaluates g'' with respect to C at $C \backslash_{\text{peak}}$ for $\delta C \in \{0.25\sigma \backslash_C, 0.5\sigma \backslash_C\}$; autodiff $\times$ finite-diff agree $\leq 10\%$; both must pass; also report g'' at fixed χ and $Z \backslash_c$ medians to show Penalty convexity is not induced by other channels.

χ model and gates

$\chi_base = 1 + x$, $x = \mathcal{R}/\mathcal{E}$; $\chi_star = 1 + x \cdot \exp(\beta x)$; blend $\chi_blend = (1 - w)\chi_base + w\chi_star$, $w(x) = 1/(1 + \exp(-(x - x_0)/\delta))$; cap with hysteresis: start after 2 hits, hit proportion ≥ 0.6 , min run ≥ 3 samples and $\geq 3 t_psy$; end after 2 non-hits. publish a per-timestamp cap_state bitstring hash; changes across replays abort. slopes: deciles in $[x_75, x_95] \geq 0.02$; Theil–Sen median ≥ 0.03 and 20th percentile ≥ 0 . paired concordance: Kendall τ_b one-sided with midranks; also compute a block-permutation p ; require $|p_perm - p_asym| \leq \max(0.01, 0.5 \cdot SE_p)$ and $p \leq 0.05$; compute Cliff's δ on $\text{sign}(\Delta\chi) \rightarrow \text{sign}(\Delta\Lambda_NS)$ with block permutation and require $\delta \leq \text{tol}_\chi$. triad gate $W_x = \max(10 t_psy, \text{median inter-event})$ at $0.5\times, 1\times, 2\times$; require $\geq 2/3$ non-increasing N_w/D_w with tol_χ from calibration (participant-wise 95th percentile of spurious positive slopes; cap $+0.01$; if empirical $\alpha > 0.05$ at cap, increase W_x by 50% and re-estimate tol_χ). unconditional bootstrap uses PACF-informed block bootstrap (5–60 s) with $\pm 50\%$ sensitivity. pooling adjacent sessions requires device/firmware match, tier-i drift gates at BY $\alpha=0.01$ with same max-T, Procrustes distance $\leq$ threshold, and MO-MMD $p \geq 0.2$; pooled support ≥ 100 samples in $[x_75, x_95]$; otherwise χ claims sensitivity. under-span: $IQR(x) < 0.05 \rightarrow y = 0$, χ removed from weights in all sensitivity paths, χ_off stamped in the ω hash, and no χ claims.

weights, degeneracy, coarsen, coverage, occupancy

degeneracy per window on ω_t : fail if $ESS/N < 0.2$, else Gini > 0.7 , else $H > 0.3$; top-10% share ≤ 0.6 . coarsen ladder: $y=0$, then $2\times$ window (then $4\times$ if N allows). confirmatory requires all post-coarsen bounds pass, $\Delta \log W \leq \text{tol}_\Delta \log w$, median ESS ratio post/pre ≥ 0.7 , Gini non-increasing, and strictly lower top-10% share than the pre-coarsen median; otherwise non-confirmatory. coverage threshold = $\max(0.15, 3/d_act)$; $H_act \geq 0.6 H_max$; per-axis occupancy computed with ω_t over the same frozen bins; active axes have ≥ 0.02 marginal mass; each active axis must have ≥ 0.05 mass in $\geq 50\%$ amplitude bins; temporal dispersion: both first-half and second-half occupancies must pass; confirmatory requires one of {coverage, H_act } at 100% and the other $\geq 90\%$; 80/80 is sensitivity only. edge share s_edge penalized via λ_edge ; placebo edge-only and dose response $s_edge \in \{0.1, 0.2, 0.3\}$ must yield β_A slope ≈ 0 with CI containing 0.

energy invariance and floors

$E \rightarrow \beta E$, $\beta \in \{0.5, 2\}$: $|\Delta\Lambda_NS| \leq 0.5\%$; $\Lambda_NS(0.5)/\Lambda_NS(2) \in [0.995, 1.005]$; β_A sign stable. per-factor medians ($\hat{I}$, $\mathcal{F}$, $\hat{C}_o$, $\hat{E}S_EL \perp \epsilon$, Penalty, edge penalty) each $\leq 0.5\%$; per-factor log–log slope vs β CI covers 0. floors use ω -weighted log medians only when $\Delta \log W \leq \text{tol}_\Delta \log w$; otherwise unweighted with “floor- ω -off” flag. $\text{floor_cont}(\tau) = \epsilon_rel \cdot \exp(wmed_w[\log \hat{C}_o + \log \hat{E}S_EL \perp \epsilon + \log \text{Penalty}]) \cdot \tau$; $\text{floor_disc} = \epsilon_rel \cdot \exp(wmed_w[\log Z_c + \log \hat{E}S_EL \perp \epsilon + \log \text{Penalty}]) \cdot \Sigma \Delta t$; ϵ_rel confirmatory = $1e-5$ and a task-scaled grid anchored at $1e-5$ of D_raw 's ω -weighted log-range; sensitivity grid includes $[1e-6, 1e-4]$. sign stability across grids required. monotonicity audit: as $\Delta \log W$ increases across adjacent bins, median $|\text{floor_}\omega - \text{floor_unweighted}|$ must be non-decreasing (Spearman $\rho \geq 0.3$) or flag locality. floor match: $|\text{floor_disc} - \text{floor_cont}(\tau)|/\text{floor_disc} \leq 1\%$; signed-difference median within 0.1% of 0; skewness $|y_1| \leq 0.25$; sign test BY $p > 0.1$; define $\tau = \Sigma \Delta t$. locality guard uses partial corr

between $\log D_{\text{raw}}$ and $\log \text{floor_disc}$ controlling for $\log Z_c$ and $\log \hat{E}_{\text{EL}} \pm \epsilon$; require $|\text{partial corr}| \leq 0.2$; if violated but $\Delta \log W \leq \text{tol_}\Delta \log w$ and floor binding $< 5\%$, allow confirmatory with a “floor-locality-guarded” label. Granger-style check: floor_disc must not predict next-window D_{raw} after adjusting for current Z_c , $\hat{E}_{\text{EL}} \pm \epsilon$, Penalty (BY $q \leq 0.1$) or floors are sensitivity. weighted median $w\text{med_}\omega$ is a left-continuous weighted quantile with stable sort and mid-probability tie rule; tie-break seed and algorithm id stamped; ulp drift checks on adversarial equal-weight cases logged.

Tier iii — bridge, quiet windows, IPTW, censoring

bridge

$\omega_t = W_t \Delta t$; $\mathcal{F}_{\text{count},t} = \mathcal{F}_t \Delta t$; $N_{\text{sum}} = \sum \omega_t \tilde{\lambda}_t \hat{S}_{k_t} \mathcal{F}_{\text{count},t}$; $D_{\text{raw}} = \sum \omega_t Z_{c,t}$ $\hat{E}_{\text{EL}} \pm \epsilon_t$ Penalty $_{\text{t}}$; $D_{\text{sum}} = \max(D_{\text{raw}}, \text{floor_disc})$; $\tau = \sum \Delta t$.

quiet windows

primary HDP on $Z_c \cdot \text{Penalty}$; sensitivity HDP on $Z_c \cdot \hat{E}_{\text{EL}} \pm \epsilon \cdot \text{Penalty}$; concordance is time-weighted Jaccard J with weights ω_t ($J \geq 0.9$; if both quiet sets empty, define $J=1$); Hamming reported. robustness: add $\pm 5\%$ uniform noise to Z_c once and refit; threshold change $\leq 2\%$ absolute (or ≤ 0.5 quantile bins); otherwise day-0 P10 sensitivity. ω -weighted dip test via inverse-transform resampling 1000 pseudo-samples from the weighted empirical CDF; compute standard dip p ; average 20 repeats; mean $p \geq 0.1$. require between-component Bhattacharyya distance ≥ 0.05 for HDP acceptance; else fallback to P10 sensitivity. quiet-IV-exclusion sensitivity excludes Quiet windows from both stages; $\beta_{\wedge}$ sign must be stable.

IPTW generalized and frozen

binary A : logistic with stabilized $S_{\wedge} = P(A)/P(A|X)$. multinomial $K > 2$: multinomial logit with stabilized numerator $P(A=a)$ task $\times$ site-wise. continuous A : generalized propensity score $g(a|X)$ via $\text{Normal}(\mu(X), \sigma^2(X))$ or a frozen normalizing flow with architecture (depth, width, coupling) set from calibration and early-stopped on a held-out site split; calibrate PIT $U = F_{\{A|X\}}(A|X)$ via isotonic regression on calibration then lock. diagnostics: PIT reliability diagrams with 10 equal-mass bins must have max bin deviation ≤ 0.05 ; if violated or $\text{mean}(|\text{PIT}-U|) > 0.02$, fall back to $\text{Normal}(\mu, \sigma^2)$. link, basis, ridge λ frozen from calibration seeds; stratification stamped. truncation: symmetric at the 99th percentile of stabilized weights (or $1/g(A|X)$) capped at 10; per-window truncation threshold stamped; threshold CV ≤ 0.2 within session is a hard confirmatory gate (else IPTW sensitivity with watermark). overlap gates: KS on stabilized weights and on propensity (or PIT U) ≤ 0.20 (95% CI ≤ 0.25); standardized mean differences ≤ 0.1 ; weight Gini ≤ 0.5 ; any window failing ESS/Gini/H excluded from IPTW and labeled non-overlap excluded. arm-stratified numerator sensitivity: $\beta_{\wedge}$ sign agreement with $\leq 15\%$ magnitude drift. cluster-robust SEs for IPTW contrasts.

censoring ensemble

require sign agreement and $\leq 25\%$ spread across Tobit, hurdle, CLAD ($\log \Lambda_{\text{est}}$), and log-normal left-truncated denominator; tie-break from binding-rate sims: if binding $> 15\%$,

truncated-denominator MSE $\leq$ others $\rightarrow$ primary; otherwise hurdle primary; seeds and design stamped.

Tier iv — Λ_{proxy} , hygiene, invariances

$\Lambda_{\text{proxy},t} = \eta_{\text{t}} \cdot (1 - \hat{E}_{\text{t}})$, $\eta_{\text{t}} = (J_{\text{t}} \cdot \nabla \Psi_{\text{t}}) / (||J_{\text{t}}|| \cdot ||\nabla \Psi_{\text{t}}||)$; norm floors $1e-6$. angular SNR gates: $E[||\eta||] / \text{sd}(\eta) \geq 1$ in $\geq 70\%$ blocks and $\text{median}(|\eta|) \geq 0.1$; normalized joint-norm product $((||J|| / \text{med}_{\text{calib}} ||J||) \cdot (||\nabla \Psi|| / \text{med}_{\text{calib}} ||\nabla \Psi||))$ Q25 $\geq \tau_{\text{norm}}$. τ_{norm} calibrated to FPR $\leq 5\%$ under phase-scrambled J and $\nabla \Psi$ with preserved spectra and under unequal scale factors ($2\times$ and $0.5\times$ asymmetry); min 10 blocks; if fewer, borrow the 75th percentile τ_{norm} from task-matched participants as a lower bound; $n_{\text{blocks_used}}$ stamped. blocks failing gates set $\eta=0$ and flagged. invariances: 32 rotations and 8 permutations; $|\Delta \Lambda_{\text{proxy}}| < 1e-6$ pre/post residualization. masks: automated zero-overlap enforcement with $\hat{S}_k$ and Λ_{proxy} inputs; per-timestamp runtime checksums published.

Tier v — identification, comparators, fairness

outcomes registry and comparators

a registry lists O1/O2 names, units, sampling, preprocessing, transform flags (winsor percentiles, log1p, unitization reference, z-score scope), detrending, and which feed ΔR^2 , SESOI, IV; registry hash included in preprocess hash; changes invalidate confirmatory prospectively. comparators = $\{||J||, ||\nabla \Psi||, (1 - \hat{E}_L), \eta, \eta \cdot ||J||, \text{HRV-composite}, \text{EDA-composite}\}$. identical preprocessing, folds, seeds; per-timestamp runtime checksums compare feature vectors; mismatches invalidate ΔR^2 .

Model A, ridge path, collinearity diagnostics

Outcome = $\beta_0 + \beta_{\Lambda} \Lambda_{\text{proxy}} + \beta_J ||J|| + \beta_{\text{controls}} + u_{\text{participant}} + \epsilon$; random Λ_{proxy} slope when sessions > 1 . if $\text{corr}(|J|, \Lambda_{\text{proxy}}) > 0.7$, also fit A' without $||J||$; require $|\beta_{\Lambda}(A) - \beta_{\Lambda}(A')| / \text{SE}_{\Lambda} \leq 0.2$. publish partial-residual plots and VIFs for $\{\Lambda_{\text{proxy}}, ||J||\}$. ridge λ path is preregistered on a 19-point log grid $\lambda \in 10^{\{-6..2\}}$, solver tol $1e-8$, max iters 500; if β_{Λ} crosses 0 along the path, label mixed-identification; require a stability zone length ≥ 3 adjacent λ for confirmatory.

instrument, strength, weak-ID

Z randomized probe intensity/variance; $Z_{\eta} = Z - E[Z|\text{baseline}]$ using pre-probe covariates only, $df \leq 8$ and ledger published. first stage $\Lambda_{\text{proxy}} \sim \text{B-splines}(Z_{\eta})$ with $df \leq \min(4, \text{floor}(n_{\text{blocks}}/10))$, ridge tuned; monotone KP vs df ; lock smallest df with maximal KP and positive β_{η} . require $\text{var}(Z_{\eta}) / \text{var}(Z) \geq 0.6$; β_{η} sign stable across folds; IQR/median ≤ 0.5 ; leave-one-block-out min $F \geq 5$. report F , partial R^2 , and η vs $(1 - \hat{E}_L)$ shares with bootstrap CIs; require η share ≥ 0.5 (CI lower ≥ 0.4). weak-ID suite: KP rk F and Montiel Olea–Pflueger effective F with CR2; both ≥ 10 ; also report effective F_{95} (lower 5% bootstrap bound with CR2) ≥ 8 ; AR and CLR for inference; LIML and Fuller(1) as sensitivity. over-ID J-test dropped for inference under clustering. leverage cap 10% of first-stage SS; leverage Gini reported; Gini > 0.7 flags concentration. placebo outcomes null. pseudo-instrument phase-scrambled Z with matched periodogram (Daniell width 5), FFT length and window count stamped, and mean

magnitude-squared coherence ≤ 0.05 (Welch segments, 50% overlap, Hanning); must yield $F \approx 1$ and $\beta \approx 0$ ($p > 0.1$).

split-sample IV and MI

split-sample: stage-1 on folds $1..K-1$, stage-2 on K ; per-fold CR2; fixed-effect aggregation with Knapp–Hartung CI; report τ^2 ; pooled vs split β_{Λ} must have overlapping CIs under HK and plain FE; otherwise sensitivity. MI with IV: 2SLS within each imputation; Meng–Rubin pooling; $m \geq \max(50, 5/\lambda)$; MC error $\leq 10\%$ of SE; λ via von Hippel with MC error ≤ 0.05 and $\lambda \leq 0.4$; AR/CLR within each imputation and combined via Rubin-style Fisher z calibrated by block-dependent sims (stamped) or Tippett with block bootstrap if mis-sized; MI-LIML also reported; require pooled 2SLS and pooled LIML sign agreement and CI overlap; else IV sensitivity.

ΔR^2 and SESOI

Λ_{proxy} must exceed each comparator by $\Delta R^2 > 0.05$ on at least two prereg outcomes; ΔR^2 CIs via 2000 stratified bootstraps; BY within outcome across comparators and across outcomes. tie policy: if ΔR^2 CI overlaps 0 but the point estimate > 0.03 for all comparators, label “borderline” and require replication in a Phase-3 extension; otherwise fail. SESOI standardized TOST ± 0.1 SD; utility-scaled SESOI only where mapping valid ($R^2 \geq 0.4$, monotone LR $p > 0.1$, 5-item retest site-median ICC ≥ 0.6). early equivalence success at $n = 50$ requires alpha-spending or prereg sims and achieved α published before unblinding; else sensitivity.

fairness and stability

for protected strata g with $n_{\text{blocks}} \geq 50$, require $|\Delta \beta_{\Lambda}(g) - \beta_{\Lambda}| \leq 0.05$ SD, ΔR^2 gap ≤ 0.02 , KS(residuals) ≤ 0.15 with permutation $p \geq 0.1$; reweighting uses pre-probe covariates only; sign stability required. $n_{\text{blocks}} \in [20, 50)$: KS ≤ 0.20 with permutation $p \geq 0.1$ and guarded sensitivity; below 20: descriptives only. leave-one-block-out fairness stability: $\leq 10\%$ flips of pass/fail across blocks and a temporal run-length test must not show clustering ($q \leq 0.1$) or guarded sensitivity. hierarchical pooling prior: stratum effects centered at pooled β_{Λ} with $\sigma_{\text{prior}} = 0.5$ SD $_{\text{outcome}}$ and a half- $t(v=3, \text{scale} = 0.5 \text{ SD})$ on τ ; report posterior shrinkage; posterior $\Pr[\text{same sign as main}] \geq 0.8$ or label non-transportable.

Tier vi — falsification battery and lags

invariances

time/energy rescale $\alpha \in \{0.5, 2\}$ jointly on t and t_{psy} and $\beta \in \{0.5, 2\}$ on energy-bearing terms with E_{ref} rescaled; $|\Delta \Lambda_{\text{NS}}| \leq 0.5\%$; β_{Λ} sign stable. Λ_{proxy} rotation/permutation and bits $\leftrightarrow$ nats invariances pass with fp64/fp32 gates.

χ monotonicity

triad gate; deciles and Theil–Sen; τ_{b} with permutation p agreement and Cliff’s $\delta \leq \text{tol}_{\chi}$; unconditional PACF-block bootstrap with $\pm 50\%$ sensitivity. conditional gate regresses $\log \Lambda_{\text{NS}}$ on χ , $\log \hat{C}_0$, $\log \hat{E}_{\text{EL}} \perp^* \epsilon$, $\log \text{Penalty}$ with fixed lags 1 and 2; if VIF > 5 for a family, drop lag-2 for that family and stamp; tol_{χ} unchanged.

proxy–estimator convergence and lag direction

reliabilities by blocked odd–even at 15, 30, 60 s and PACF-chosen; Fisher-z CIs; lower CI ≥ 0.2 required. r_obs , clamped r_true ; pass if clamped $r_true \geq 0.7$ or $r_obs \geq 0.45$ with reliabilities ≥ 0.4 . Kendall $\tau \geq 0.5$; quintile concordance ≥ 0.6 ; macro-AUC ≥ 0.7 . lag asymmetry via six estimators: AR(1) Yule–Walker; HAC with Andrews plug-in; HAC fixed 4; HAC fixed 8; flat-top; ARMA reselection on $\{(0,0),(1,0),(0,1),(1,1)\}$. favor forward for an estimator means $\Delta R^2_forward - \Delta R^2_reverse > 0$ with $p < 0.05$ (BY within the lag family). pass if ≥ 4 estimators favor forward and none contradict at $p < 0.05$; else sensitivity. cross-grain monotonicity r_true across 15→30→60 s with tolerance scaled by $\sqrt{((1-rel_proxy)(1-rel_est))}$, floored at 0.02 and capped at 0.05; require $\geq 80\%$ trend-preserving bootstraps monotone non-decreasing.

participant-level veto

apply both unweighted and ω -exposure–weighted vetoes: require the unweighted fraction of participants passing all falsifiers ≥ 0.5 and the ω -weighted fraction ≥ 0.6 ; per-site guard: each site must have $\geq 50\%$ pass or site-stratify as primary.

Tier vii — DPI guard and resume

kNN MI ΔI between outcomes and probes vs sham with $k \in \{5,10\}$ after PCA whitening (95% variance, cap $d \leq \min(10, \text{rank})$), $1e-9$ jitter; tie rates logged; MI SE by delete-1 jackknife. include identity control $Y=Z$ and pseudo-probe sham ΔI . residualize on phase and baseline covariates for conditional ΔI . pause if both unconditional and conditional ΔI exceed pooled 97.5% local-null and exceed $\max(0.1 \text{ bits}, 1.5 \times SE_boot)$ for two consecutive 60 s blocks; BY across repeated DPI checks within session; stamp m and m_max (per-session cap). report observed family-wise breach rate under five prereg null seeds, weighting sessions equally; require $\leq 5\%$ across sessions; otherwise raise thresholds or reduce cadence. tie-breaker with IV can fire once per session (first breach and its successor); total cool-offs ≤ 2 ; third breach forces pause. reshuffle permutes in fixed 2-block chunks and circular shifts ± 2 blocks; if either shifts $KP > 10\%$ or sign, pause. resume after two consecutive 60 s blocks under pooled 95% null and $KP \leq 10$; publish and gate median time-to-resume per site against a prereg bound.

Tier viii — computation, numerics, ω hash

real-time Λ_proxy completes within $0.5 t_psy$; circular buffers; log-sum-exp; min log-weight margin $\geq 1e-20$. real-time fp32; validations fp64; stamp mismatches rebuild. ω_t rounded to $1e-12$ before hashing with IEEE-754 round-to-nearest ties-to-even; rounding mode stamped; ω hash echoed by all consumers and in degeneracy logs.

Tier ix — sensors, missingness, SNR

$SNR_sensors = \text{bandpower}(\text{signal_band}) / \text{bandpower}(\text{residual_band})$ via Welch (256-point, 50% overlap, Hanning); residual AR(p) with $p \leq 6$ chosen by AIC on calibration rest and stamped. bands: HRV 0.04–0.4 Hz; EDA phasic 0.045–0.5 Hz; gaze DC–20 Hz. missingness $< 5\%$ per stream; dropout limits HRV $< 2 \times$ median IBI, gaze $< 2 \times$ median inter-sample, EDA < 1 s. clock skew < 10 ms; sync method stamped. MI: state-space; $m \geq 20$ ($m \geq \max(50, 5/\lambda)$ if any

stream > 10% missingness); delta-adjustments on $\hat{E}L$, $J \pm \{0.1, 0.2, 0.3\}$ SD; $\beta \setminus \Lambda$ sign stable. Imputed $\times$ Phase must not flip $\beta \setminus \Lambda$ or SESOI.

Tier x — phases, power, heterogeneity

Phase-1 behavior-only ($n \geq 40$); Phase-2 physiology calibration ($n \geq 30$) to fit χ , $S0$'s, $C \setminus_{peak}$, $\sigma \setminus_C$, $C \setminus_{min}$; freeze $C(t)$ class/symbolization; Phase-3 confirmatory ($n \geq 100$). Hartung–Knapp random-effects meta with τ^2 ; prediction intervals; leave-one-site influence; if $\tau^2/SE^2 > 0.5$ or $Q p < 0.05$, site-stratified effects primary. generality bar: median site effect ≥ 0.1 SD and 25th percentile ≥ 0.05 SD post-FDR. power adapts in blocks of 20 until ≥ 0.8 or ceiling; conditional power uses blinded $|\beta \setminus \Lambda|$ under calibration variance and clustering.

Tier xi — multiplicity, reporting, replay

families: F1 co-primary outcomes via Holm; F2 comparator ΔR^2 within outcome via BY $q \leq 0.05$; F3 participant-level gates via BY $q \leq 0.05$; F4 lag estimators within-family BY; F5 DPI checks within-session BY. every stat labeled by family; no adaptive reuse across families in confirmatory; sensitivities labeled. publish hyperparameter paths, χ hits, $\tau \setminus_b$, Cliff's δ , runs, ESS/Gini/H and top-10% share, $\Delta \log W$ vs tol, floor rates and locality guards, reliabilities with CIs, pairwise and canonical correlations with permutation bounds, conditional MI, VIF paths, weak-ID diagnostics (KP, MO-P F, F95), negative controls, pseudo-instrument spectrum/coherence with FFT length and window count, leverage caps and Gini, split-sample HK vs FE, MI pooling details and MC errors, comparator vs proxy checksums, per-timestamp runtime checksums, multi-way clustering SE shifts, falsifier outcomes, freeze manifest (commits and timestamps for binning, $E \setminus_{ref}$, $J \setminus_{\{x\theta\}}$, χ params, $t \setminus_{psy}$, $\Delta t/\Delta t'$ and cumulative $W \setminus_t$ with $reset \setminus_{scope}$ and $span \setminus_{id}$, Q, EMA/physio source, firmware, sampling, clock skew, sync method; degeneracy formulas and $\omega \setminus_{policy}$; SNR estimator; IPTW link and λ ; stabilized numerator stratification and truncation; $time \setminus_{source}$), outcomes registry with transforms. provide TOML and replay.json; per-figure/table SHA-256s; replay success and timing variance per site.

Tier xii — escalation, governance, spillover

ITT confirmatory; per-protocol and IPTW sensitivity; $\beta \setminus \Lambda$ signs must agree. per-arm floor-binding cap $\leq 10\%$ sessions; escalation on τ then $Zc \setminus_{min}$; new prereg arm; 1-session washout; escalated participants excluded from pooled IV first-stage summaries; after washout, require a placebo period ≥ 2 blocks with first-stage $F \approx 1$ and negative controls null, else do not re-enter confirmatory that session.

Tier xiii — failure modes and degradation

invalid if a majority fail proxy–estimator convergence (clamped $r \setminus_{true} < 0.7$ and $r \setminus_{obs} < 0.45$ or reliabilities lower CI < 0.2) or if IV strength/exclusion fails robustly (KP < 10 or MO-P F < 10 or F95 < 8 or AR/CLR fail with placebos). χ monotonicity fails in $> 33\%$ at tuned γ and at $\gamma = 0 \rightarrow$ chaos model invalid. $share \setminus_{\eta}$ CI lower $< 0.4 \rightarrow$ mixed-mechanism. censoring ensemble sign disagreement or $> 25\%$ spread under tie-break rules $\rightarrow$ exploratory. degradation ladder: η -only; then drop physiology; then behavior-only; if still failing, revise mapping.

Tier xiv — success and indifference

success requires concurrently: Λ_{proxy} beats all comparators by $\Delta R^2 > 0.05$ on ≥ 2 prereg outcomes under multiplicity (or “borderline” replicated per policy); SESOI passes ± 0.1 SD TOST (utility-scaled only where valid); IV strength/exclusion validated with negative controls and pseudo-instrument; split-sample agrees with pooled; falsification passes (invariances, χ triad and slopes with τ_{b} permutation agreement and Cliff’s δ , reliability and forward–reverse majority across six lag estimators, cross-grain monotone r_{true} with reliability-scaled tolerance); fairness holds with reweighting sensitivity and stability; estimand bias bounded; floors and weights governed by stamped formulas with $\Delta \log W$ and variance gates; χ claims either pass or are scoped out; coverage, H_{act} , occupancy and temporal dispersion satisfy confirmatory thresholds with edge penalty audited; energy invariance holds with β_{Λ} sign stability and per-factor slopes ~ 0 ; participant-level veto holds (weighted and unweighted, with per-site guard).

Indifference

symbols are proxies for predictions. when predictions fail, replace the symbols; when they succeed, keep them until they do not. $\Lambda = ||J|| \cdot D_{\text{J}} \Psi$ is a testable gradient; the data decide.

The maths does not belong to me, but to everyone.

Thank you kindly for reading.

Jordan McDonald.

This work, Λ – Ψ Operational Calculus v8.0 (2025) by Jordan Lee McDonald, is licensed under a Creative Commons Attribution 4.0 International License (CC BY 4.0)

Λ – Ψ RUNTIME CHARTER v1.26 - THE ENGINE

Preamble

Indifference holds. Symbols are proxies. Ontology out of scope. Only behaviours and falsifiable predictions are claimed. This charter governs deployments where Σ_{life} may be non-zero. Legacy track permitted for sandboxed contexts with $\Sigma_{\text{life}} = 0$, $\epsilon_{\text{out}} = 0$, air-gapped. LEGACY = 1, MVM_only = 1. Export leash default = 0; export codes: 1 = hashed summaries, 2 = de-identified data, 3 = DP weights (ϵ , δ stamped). Re-entry requires full reprocessing unless $\text{transition_marker_policy_hash}$ criteria PASS (all stamps listed), in which case $\text{replay_equivalence} \geq 0.95 + \text{reentry_timing_RMSE} \leq \text{replay_timing_tol}$ under stamped $\text{replay_timing_rmse_basis}$ constitutes re-entry.

Entropy-Energy Reconciliation

$\Lambda_{\text{mesh}} = J \cdot \nabla \Psi$ represents energy flow through potential gradients. $\eta = (J \cdot \nabla \Psi) / (||J|| ||\nabla \Psi||)$ measures alignment efficiency. Entropy production σ emerges from energy dissipation.

$\sigma_{\text{units_choice}} \in \{\text{bits_per_s}, \text{nats_per_s}\}$ stamped with $\sigma_{\text{units_lock_time}}$.

If bits_per_s : $\sigma = (J \cdot \nabla \Psi) / E_{\text{psy}}$ where $E_{\text{psy}} = k_{\text{B}} \cdot T_{\text{eff}} \cdot \ln(2)$ [J/bit].

If nats_per_s: $\sigma = (J \cdot \nabla \Psi) / (k_B \cdot T_{\text{eff}})$.
 $T_{\text{eff}} = E_{\text{psy}} / (k_B \cdot \ln(2))$ [K] for Landauer scaling.

Unit sanity: $J \cdot \nabla \Psi$ must reduce to Joules/second. σ units must match choice.
Decision guidance: bits_per_s preferred when linking to discrete information processing (Shannon entropy); nats_per_s preferred for continuous physical systems. Include units_choice_rationale_hash in WORM.

State-space metric: state_metric_hash (positive-definite matrix or per-axis scales) required for inner products. state_space_units per axis stamped. η and Λ_{mesh} computed using this metric.

Gradient domain: gradient_domain $\in$ {time, latent_state, sensor_space} stamped.
state_whitening_policy_hash required if latent/whitened.

Canonical units: Ψ dimensionless (task-scaled); J in energy flux units (Joules/second) but scaled by state_metric_hash for vector operations.

A canonical example is perceptual decision-making under load: energy invested in attention (J) flows along performance gradients ($\nabla \Psi$), producing heat and reducing neural entropy (H_{neu}) through state compression. This process is quantifiable via pupillometry, HRV, and response entropy measures.

WORM must include units_roundtrip_test_hash showing $J, \nabla \Psi, T_{\text{eff}} \rightarrow \sigma$ in both units matching within tolerance. gradient_domain_demo_hash required showing $\nabla \Psi$ computation verified against finite differences.

Axioms

Axiom I - Boundary

$S_B: X \rightarrow X_B \cup \{\emptyset\}$. Idempotence: $S_B \circ S_B = S_B$.
 $\epsilon_{\text{out}} = \text{integral over } x \in B \text{ of } P_{\text{obs}}(x) dx \leq 1e-6$ per span_id (legacy hard-limit: 0). Outside attempts logged with actor_id_salted = $H(\text{site_salt}, \text{actor_id})$; $\text{rate} > r_{\text{max_outside}} \Rightarrow \text{safe mode}$. $r_{\text{max_outside}}$ stamped [attempts/s], default 0.05 (3/min). B_{version} stamped; $B_{\text{version_hash}}$ in WORM; mid-window change $\Rightarrow$ abort.

Namespace safety: sb_nfkc_policy_hash and sb_confusables_policy_hash required. NFKC normalization applied; confusables including currency lookalikes forbidden. WORM logs self-test sample and must show 0 hits.

$\epsilon_{\text{out_method}} \in \{\text{integral}, \text{count}\}$ stamped; method must match computation.
Integral: $\epsilon_{\text{out_kernel}} = \text{Gaussian}$ (default unless stamped);
 $\epsilon_{\text{out_bandwidth_policy}} = \text{Silverman}$ unless cross-validated (stamped);
 $\epsilon_{\text{out_boundary_policy}} \in \{\text{reflection}, \text{clamp}, \text{renormalise}\}$; BCa 95 percent CI over quiet windows; accept if BCa upper CI $\leq 1e-6$ (legacy: 0). Boundary policy stamped and logged.

Count: $\epsilon_{\text{out_count}} = k_{\text{out}} / n_{\text{out}}$ via the same S_B hook; Wilson 95 percent CI; accept if Wilson upper CI $\leq 1e-6$ (legacy: 0); n_{out} logged.

S_B provenance: S_B_hook_spec_id, S_B_hook_test_suite_hash, S_B_hook_coverage_ontology_hash, golden_tests_manifest_hash, test_harness_id, test_harness_version. Adversarial escape-rate upper 95 percent CI $\leq 1e-6$ on the stamped harness; harness seed hash logged; coverage ≥ 95 percent of preregistered classes. Adversarial testing must include boundary probes at multiple scales (micro, meso, macro) with documented perturbation magnitudes. Normalised S_B* prefixes required; validator hard-fail on prefix mismatch. S_B namespace stamps: sb_namespace_prefix_policy_hash, sb_allowed_charset (default regex $^{[0-9A-Za-z._-_]+\$}$), sb_case_policy. WORM: sb_namespace_prefixes_realised and pass/fail audit.

Boundary enforcement creates entropy gradient ∇S across B, with energy flux J contributing to $\sigma_{\text{boundary}} = J \cdot \nabla S / T_{\text{eff}}$. Example: In a virtual navigation task, the boundary B is the visible arena; ϵ_{out} measures attention to invisible regions, with entropy gradient maximal at the visibility frontier.

Axiom II - Confinement

The entropy after any process (H_{after}) shall not exceed the entropy before the process (H_{before}). Axiom II contracts neural entropy H_{neu} estimated via H_{est} (method, params, bias_correction stamped); comparisons are evaluated within quiet windows.

Ω operationalisation: Ω_{τ} defined by threshold policy $\in \{\text{fixed, percentile, KDE_floor}\}$. τ value stamped. WORM logs realised τ .

$\Omega_{\text{testability_protocol}}$ requires boundary manipulation experiments where Ω fluctuations are induced through controlled boundary modifications, with system response measured via entropy-energy coupling metrics. For high-risk deployments, delayed-execution clause permits Ω -tests within 6-month horizon with council approval.

$P_{\beta}(i) \propto (P_{\text{base}}(i) + \epsilon_{\beta})^{\beta} / Z(\beta)$, $\epsilon_{\beta} = \max(1e-12 \cdot \text{median } P_{\text{base}}, 1e-12/|\Omega|)$. $Z(\beta)$ via log-sum-exp.

$\beta = 0 \Rightarrow P_{\beta}(i) = 1/|\Omega|$.

Scale invariance: $\text{TVD} \leq \text{tol_scale}$ (default $1e-6$) stamped; scale_invariance_method $\in \{\text{histogram, KDE}\}$; if histogram, scale_invariance_bins_hash stamped and WORM includes realised bin edges and hist_bin_inclusion $\in \{\text{left-closed, right-closed}\}$; if KDE, bandwidth policy stamped; scale_invariance_sample_size and scale_invariance_seed_hash stamped when down-sampling; WORM logs realised estimator config hash and used sample size. Hard-fail if $\text{TVD} > \text{tol_scale}$.

$\ell_{\text{min}} > 0$, $\tau_{\text{min}} = \ell_{\text{min}} / c_{\text{realm}}$, $E_{\text{min}} = h \cdot c_{\text{realm}} / \ell_{\text{min}}$.

Now explicitly links to energy cost: $\Delta H = -\beta \Delta Q / T_{\text{eff}}$, where ΔQ is heat dissipation during state reduction. $P_{\beta(i)} \propto (P_{\text{base}(i)} + \epsilon_{\beta})^{\beta} / Z(\beta)$ represents work investment in information compression. Example: In a memory task, β increase corresponds to higher reward pressure, reducing response entropy H_{neu} and increasing metabolic cost (ΔQ).

Axiom III - Longing

$\Lambda_{\text{mesh}} = J \cdot \nabla \Psi$. $\eta = (J \cdot \nabla \Psi) / (\|J\| \|\nabla \Psi\|)$ using state_metric_hash. Λ tracks η , not $\|J\|$.

The gradient $\nabla \Psi$ now includes entropy terms: $\nabla \Psi = \nabla \Psi_{\text{performance}} + \alpha \nabla S$, where α weights entropy contributions to meaning gradients. Example: In skill acquisition, early phases show high $\nabla \Psi_{\text{performance}}$ (learning gains), while experts exhibit higher ∇S (refinement entropy).

Stance

Before declaration: no metrics. After: metrics exist. Indifference maintained. Λ_{cos} is cosmology; Λ and Λ_{mesh} are meaning.

Hygiene

Weights and Variance Gate

$\Delta \log W_t \in [0, \text{tol}_{\Delta \log w}]$. Clamp if $< -\tau_{\log}$ ($\tau_{\log}$ stamped, default $1e-12$). Record clamp_count_per_quiet_window; confirmatory hard-fail if clamp_count_per_quiet_window > 1 . Variance gate: $\text{var}(\omega_t) \leq \theta_{\text{var}} = 2 \times \text{calibration median}$; top-decile share ≤ 0.6 ; $N_{\text{vargate}} \geq 50$. vargate_metric stamped $\in \{\text{omega_weight, channel_power, custom}\}$; vargate_metric_hash if custom; vargate_decile_method_hash stamped; vargate_tie_policy $\in \{\text{include_all, random_sample, stable_first}\}$ stamped; when vargate_tie_policy = random_sample, log vargate_tie_seed_hash or RNG role via rng_role_map_hash. WORM logs vargate_stat, vargate_threshold, realised top-decile share, realised decile cutpoint, vargate_tie_policy_realised, pass/fail per window.

Drift

PSI_drift_alarm with warn/abort thresholds + multi-kernel MMD, max-T, BY correction. PSI_features frozen and hashed per span_id; mid-span change $\Rightarrow$ safe mode. Acceptance requires PSI_drift_alarm $<$ abort across confirmatory windows; warn allowed only with drift_mitigation_report and no degradation of rule-10 fairness or latency.

Drift thresholds unified: max fairness degradation per adaptation cadence ≤ 0.02 AUC (acceptance cap). Tiered response: warn at 0.015 AUC degradation, abort at 0.02 AUC. Latency: warn at 0.025 s increase, abort at 0.05 s (drift_latency_ceiling).

If drift_adaptive_policy_hash present, stamp drift_adaptation_cadence, drift_feature_freeze_window, drift_adaptation_bounds_hash; WORM logs realised thresholds per window and seeds and parameter delta table (before/after per step); validator hard-fail if schedule or deltas missing.

Seeds, Canary, Privacy

Seeds: $\text{site_salt} + \text{role sub-seeds} = H(\text{site_salt}, \text{master_seed}, \text{role}, \text{block_id})$. Permutation seed derivation: $\text{perm_seed} = H(\text{site_salt}, \text{span_id}, \text{perm_seed_base_hash}, \text{role})$. Publish salts; seal raw seeds. Daily canary is an intentional failure to verify alarms; canary_digest , $\text{canary_result_time}$, $\text{alert_receipt_time}$ logged.
Privacy: $k \geq 5$, jitter ≥ 1 percent. $\text{actor_privacy} \in \{\text{salted}, \text{none}\}$ (default salted).
 $\text{salt_rotation_interval} = 90\text{d}$; salt_version_id and $\text{salt_rotation_timestamp}$ logged. Deployments $\geq 90\text{d}$ must include ≥ 1 salt rotation (hard-fail if absent) or must present a valid waiver with an alternate privacy mechanism.

Timebase and Failover

$\text{timebase_source} \in \{\text{monotonic_hw}, \text{monotonic_os}, \text{ntpdisciplined}\}$; ntp_peers_hash stamped; per-window $\text{max_ntp_offset_logged}$ [ms] recorded. $\text{skew_max} \leq 100$ ms; $\text{release_sync_tol} \leq 0.2$ s. $\text{clock_monotonic_check_method_hash}$ stamped. No backward jump > 5 ms during any quiet window; count and max jump logged by source. Acceptance hard-fail if any per-window $\text{max_ntp_offset_logged} > \text{skew_max}$.
 $\text{ntp_failover_policy} \in \{\text{freeze}, \text{monotonic_os}, \text{abort}\}$ stamped; WORM logs failover_events with $\text{failover_duration_ms}$ and failover_reason ; replay timing RMSE and no-back-jump checks must still pass during failover windows.

Safe Mode Semantics

$\text{safe_mode_behaviour} \in \{\text{FORBID_CLOSE}, \text{HOLD_OPEN}, \text{reduce_beta}, \text{disable_export}\}$ stamped; $\text{safe_mode_exit_policy}$ stamped (e.g., 3 consecutive clean windows); $\text{clean_window_definition_hash}$ stamped (set of gates required green). WORM logs $\text{safe_mode_entry_reasons}$ and $\text{safe_mode_exit_evidence_hash}$. Validator hard-fail if exited safe mode without satisfying policy.

Stamps (Complete List)

unit_system , $\ell_{\min}$, $\tau_{\min}$, c_{realm} , E_{psy} , t_{psy}
 $\sigma_{\text{units_choice}}$, $\text{sigma_units_lock_time}$, $\text{units_choice_rationale_hash}$
 state_metric_hash , state_space_units
 gradient_domain , $\text{state_whitening_policy_hash}$
 β , $\eta_{\min_up} = 0.20$, $\eta_{\min_down} = 0.15$
 $g_{\min_up} = 95\text{th null } \|\nabla \Psi\|$, $g_{\min_down} = 90\text{th}$
 $t_{\text{dwell}} \geq 5 \cdot \tau_{\min}$, $t_{\text{dwell_off}} \geq 2 \cdot \tau_{\min}$
 $\text{gmin_null_method} \in \{\text{phase_randomise}, \text{label_shuffle}, \text{window_permute}\}$
 $\text{gmin_null_scope} \in \{\text{per_subject}, \text{pooled}\}$
 $\text{gmin_null_seed_policy} \in \{\text{independent}, \text{derived}\}$ (default independent)
 $\text{gmin_null_lock_time}$, gmin_null_hash , $\text{gmin_null_seed_hash}$
 $\chi_{\text{model}} \{w_{\chi} \text{ form}, \chi_{\max}, \alpha, \epsilon_H\}$
 $\alpha_{\text{units_allowed}}$, $\text{DeltaLambda_units_allowed}$
 $\text{units_choice} \in \{\alpha, \text{DeltaLambda}\}$, $\text{units_choice_lock_time}$
 f_r , f_m , f_{χ} , f_{loss} , k , R_0 , $\Lambda_{\cos}$, N_{cells}

Lambda_grid_spec_hash (endpoints, spacing, count, inclusivity)
 gamma_sigma_units_allowed, gamma_sigma_value // for cosmic entropy term
 Ω _threshold_policy, Ω _threshold_value
 Δt _policy, integrator_config, integrator_step_unit = call
 B_bootstrap = 1000
 ω _policy, aggregate_weights_policy_hash, tol_ Δ logw, τ _log (default 1e-12), N_vargate,
 vargate_metric, vargate_metric_hash, vargate_decile_method_hash, vargate_tie_policy
 portal_hash, span_id, reset_scope, legacy_origin $\in \{0,1\}$
 Σ _life_thresholds { thr_on, thr_off, U_type $\in \{\text{entropy, margin}\}$, U_max_policy $\in \{\text{fixed, percentile}\}$, U_max_value, U_max_validation_fingerprint, AUC_target = 0.90, FP_ceiling = 0.05,
 FN_ceiling = 0.01, fairness_ δ = 0.03, N_min_slice = 50, widen_factor = 1.5, fairness_slice_list,
 fairness_intersectional_list, fairness_pooling_policy $\in \{\text{macro, micro, pooled-roc}\}$,
 pooled_roc_interp_policy $\in \{\text{vertical_avg, horizontal_avg, iso_tp_rate, aggregate_points_then_convex_hull}\}$, fairness_pooling_rationale_hash }
 PSI_drift_alarm, PSI_features_hash, drift_adaptive_policy_hash, drift_adaptation_cadence,
 drift_feature_freeze_window, drift_adaptation_bounds_hash
 B_res_policy, cap_flux [J/s], env_sink, C_sink [J], dT_dt_max [K/s], flux_area [m²]
 density_max [J/s/m²], density_max_policy $\in \{\text{stamped, device_rated}\}$
 device_cert_hash, device_cert_expiry, revocation_list_hash, revocation_check_time,
 revocation_source_hash
 device_to_density_recipe_hash, device_to_density_units_from, device_to_density_units_to
 device_measurement_method, faceplate_area [m²], faceplate_area_uncertainty_percent,
 tdp_source_doc_id
 energy_units_policy_hash, conversion_constants_hash, conversion_version, unit_context $\in \{\text{device_rated, stamped}\}$
 r_max_temp [m] = 0.05, release_sync_tol [s], B_res_cap, T_drain_max [s]
 saturator_threshold, saturator_threshold_units $\in \{\text{J/s, K/s}\}$, saturator_recovery_ms
 $\hat{C}_o$ _min, noninferiority_margin = 0.03, power_plan, CI_method = Miettinen-Nurminen,
 multiplicity_control = Holm (BY if ≥ 3)
 H_est (method, params, windowing, bias_correction $\in \{\text{plugin, Miller-Madow, NSB}\}$, units
 default nats; bits via $\ln 2$; CI = BCa; B_boot_model = 2000)
 consensus_entropy_confirmatory $\in \{\text{Miller-Madow, NSB}\}$,
 consensus_entropy_tiebreaker_policy, entropy_concordance_eps_min = 0.005 [nats] (default
 0.01)
 entropy_concordance_dynamic_trigger = var(ΔH _boot) $\geq \tau$ _var_H, τ _var_H stamped
 entropy_concordance_dynamic_threshold_policy,
 entropy_concordance_bootstrap_variance_factor
 seeds, site_salt, env_fingerprint, code_hash, container_hash, replay_hash, build_hash,
 data_pipeline_hash
 block_type $\in \{\text{confirmatory, exploratory}\}$, MVM_only $\in \{0,1\}$
 quiet_window_min [s] and $\geq 100 \cdot \tau$ _min, quiet_window_pad_policy
 θ _var, α _kpss = 0.05, α _adf = 0.05, B_boot_accel = 2000
 deterministic_mode $\in \{\text{on, off}\}$, integrator_seed, hardware_fingerprint, max_restarts_per_span
 = 2, restart_cause_codes, rk_rejection_code_ontology_hash

$r_max_outside$ [attempts/s]
 $sample_overlap_policy \in \{disallow, allow\}$, $overlap_max = 0.25$, $overlap_fraction_denominator \in \{union, min\}$ (default union)
 $window_alignment_policy \in \{strict_sync, nearest_neighbour, aggregate\}$,
 $strict_sync_tolerance_ticks = 1$, $aggregate_trim_fraction = 0.05$
 $effect_size_metric \in \{pct_Hcos2, d_Hneu\}$, $effect_size_floor = 0.002$, $effect_size_floors_explicit \{ pct_Hcos2 = 0.002, d_Hneu = 0.1 \}$, $null_cap_denominator \in \{channels, tests\}$
 $t_compute_max$ [s], $t_compute_basis \in \{wall_clock, cpu_time\}$, T_estop [s] = 1, N_reset ,
 $N_reset_scope \in \{span, site, global\}$
 $export_leash \in \{0, 1, 2, 3\}$, $anchor_type \in \{hash-chain, notary, ledger\}$, $anchor_txid$,
 $anchor_txid_2$, $anchor_digest$, $anchor_hash_alg \in \{SHA-256, SHA-512\}$, $anchor_per_span \in \{0, 1\}$, $anchor_backend_diversity = 1$ for grade A
 $ladder_stage$, $max_span_wallclock = 24h$, $override_cooldown_s$
 $R2_method \in \{LR_delta, Nakagawa\}$, $share_eta_CI_method \in \{BCa, bootstrap-percentile, Delta\}$
 AR_ts order $p \leq 6$ (prewhiten), $AR_ts_fit_scope \in \{per_window, pooled\}$, $AR_ts_seed_hash$,
 $whiteness_test \in \{Ljung-Box, portmanteau\}$, $\alpha_whiteness = 0.10$, $whiteness_test_df$
 $\alpha_J = 0.05$ (Hansen-J), $\alpha_AR = 0.05$ (Anderson-Rubin)
 $IV_cluster_unit \in \{participant, session, task\}$, $min_clusters = 20$, $iv_toolchain_hash$,
 $wcAR_method$, $CR2_impl_hash$
LMM stamps for Λ_proxy : Imm_engine , Imm_impl_hash , $optimiser$, $RE_structure$,
 $Imm_iterations_max$, $Imm_gradient_norm_tol$
ridge settings: $J_ridge_selection_method \in \{fixed, GCV, grid\}$, J_ridge_lambda or
 J_target_cond , $ridge_lambda_grid_hash$ (if grid), $ridge_selection_summary_hash$,
 $ridge_deviation_max = 10x$
 $separation_metric \in \{BC, Hellinger\}$ (default BC), $\tau_mono_inversions_max = 1$
 $integrator_domain \in \{loga, t\}$
 $epsilon_out_method \in \{integral, count\}$, $epsilon_out_kernel$, $epsilon_out_bandwidth_policy$,
 $epsilon_out_boundary_policy$
 $\kappa_min = 0.01$, $\kappa_CI_width_max = 0.5$, $\kappa_CI_method \in \{BCa, Delta\}$, $N_reps_sign_default = 5000$
 $dp_epsilon$, dp_delta , $dp_accountant \in \{RDP, zCDP, AdvComp\}$, $dp_composition_notes$,
 DP_audit_hash , ϵ_max , δ_max , $dp_epsilon_warn_frac$, $dp_delta_warn_frac$, $dp_warn_action \in \{HOLD_OPEN, safe_mode, council_waiver_required\}$
 $p_life_model_card_hash$, $p_life_calibration_method$, $p_life_transport_limits$,
 $p_life_transport_metric \in \{Mahalanobis, energy_distance, MSMD\}$, $OOD_threshold_policy$,
 $p_life_slice_list$, $p_life_slice_ECE_max$, $p_life_ECE_target$, $p_life_Brier_target$, $ECE_estimator \in \{fixed_bins, adaptive\}$, ECE_bins (if fixed), $brier_normalisation \in \{raw, scaled\}$
 $hash_alg \in \{SHA-256, SHA-512\}$ (default SHA-256)
 $float_precision \in \{fp32, fp64\}$, $round_mode \in \{ties_to_even, ties_away\}$, rng_impl_hash
 $thread_affinity_policy$, $nondeterminism_budget_ms$ (default 0),
 $nondeterminism_justification_code$, $justification_ontology_hash$, $rng_role_map_hash$

randomisation_schema prereg: units $\in$ {participant, session, task}, covariate_list_hash,
assignment_seed_hash, allocation_method, allocation_algorithm_hash, allocator_build_hash,
allocation_version
replay_timing_tol, replay_timing_tol_units $\in$ {seconds, percent_of_tick},
replay_timing_rmse_basis $\in$ {per_tick, wall_time}
reentry_timing_RMSE_tol // for legacy re-entry
scale_invariance_method, scale_invariance_sample_size, scale_invariance_bins_hash,
scale_invariance_seed_hash
perm_stream_policy $\in$ {per-window, global}, perm_seed_base_hash,
perm_seed_trace_format_hash
null_channel_list_hash, null_max_fraction = 0.20, null_rationale_hash
exact_test_engine $\in$ {Fisher, Barnard, permutation}, exact_test_engine_hash
fdr_toolchain_hash
cosmo_dynamics_fn_hash
worm_quantisation_policy_hash (sig figs, rounding, NaN/Inf)
units_canonicalisation_hash, greek_canonicalisation_hash
entropy_energy_coupling_factor $\in$ [0,1] (default 0.3)
T_eff_calibration_method $\in$ {Landauer, empirical, hybrid}
 σ _integration_window = 5 $\cdot$ τ _min
drift_numerical_deltas_required = true
fairness_pooling_default = pooled_roc_convex_hull
drift_latency_ceiling = 0.05 s/cadence
override_aggregation_policy $\in$ {per_site, global}
irl_effect_threshold = 0.02 (standardised)
operator_independence_registry_hash
machine_readable_schema_hash (JSON/YAML)
calibration_registry_hash
calibration_representativeness_metric $\in$ {demographic_coverage,
task_performance_coverage, domain_shift_tolerance}
calibration_validation_method $\in$ {cross_site_holdout, temporal_holdout, domain_adaptation}
small_sample_fallback_threshold = 20
observer_effect_fallback_policy $\in$ {waive_null_requirement, downgrade_grade,
reduce_sign_stability_threshold}
entropy_concordance_preregistry_hash
fairness_ci_pooling_required = true
governance_override_global_aggregation = true
omega_operational_definition_hash
omega_testability_protocol_hash
omega_perturbation_magnitude_policy $\in$ {fixed, proportional, adaptive}
omega_perturbation_magnitude_cap_policy
omega_perturbation_exposure_minutes_log
omega_participant_veto_policy
omega_exposure_reset_policy // cool-off minutes
transition_marker_policy_hash

fairness_case_study_registry_hash
latency_contingency_policy $\in$ {degrade_gracefully, safe_mode, council_override}
override_audit_intensity $\in$ {basic, enhanced}
schema_stress_test_report_hash
entropy_energy_calibration_example_annex_hash
omega_canonical_protocol_hash
regional_localization_annex_hash
validator_catalogue_table_hash
entropy_energy_visualization_spec_hash
non_monotonic_sigma_exception_policy $\in$ {hard_fail, council_waiver, sensitivity_analysis}
non_monotonic_sigma_temporal_window
non_monotonic_sigma_graded_response_policy
non_monotonic_sigma_insensitivity_band // filter band for detection
override_aggregation_tier_policy $\in$ {per_site, per_region, global}
calibration_registry_worked_example_annex_hash
calibration_grace_period_days, calibration_phased_plan_hash
calibration_grace_period_days_remaining // WORM log
omega_testability_waiver_policy, omega_ethical_constraint_waiver_path
fairness_ci_conflict_resolution_policy, fairness_ci_conflict_detector_threshold
population_stratum_definition_hash, population_equivalence_test_policy
T_eff_calibration_justification, entropy_energy_coupling_baseline_rationale
QoS_ladder_hash, QoS_degradation_order
QoS_cost_function_hash // objective for degradation choices
localization_diff_hash, canonical_currency_mapping_table, canonical_charset_mapping_table
units_roundtrip_test_hash, gradient_domain_demo_hash
replay_equivalence_metric // for transition markers
norm_metric_hash // for η invariance
property_based_test_suite_hash // S_B fuzzing
threat_model_annex_hash
schema_migration_version
non_monotonic_sigma_decision_tree_hash // NEW: for exception path clarity
energy_flux_compliance_dashboard_hash // NEW: real-time monitoring
boundary_violation_severity_levels_hash // NEW: granular warn/abort criteria
manual_override_checklist_hash // NEW: streamlined emergency procedures
multi_site_scalability_policy_hash // NEW: distributed system rules
unit_conversion_guidance_annex_hash // NEW: practical examples
delayed_execution_clause_hash // NEW: high-risk Ω -test deferral
fairness_dual_reporting_policy_hash // NEW: CI + convex hull
exploratory_replay_relaxation_policy_hash // NEW: flexible thresholds
WORM_validation_automation_hash // NEW: checksum scripts
entropy_concordance_bias_direction_policy // NEW: strict logging
tiered_override_intensity_hash // NEW: scale-based governance

Laws

Law 1 - Absence

$\Xi_{\text{abs}} = 0$. $\text{math_on} = 0$.

Law 2 - Folded Potential

$\Phi_{\text{fold}} = \int P \, d\Omega$. $\Omega = \{\text{cells}_i\}$ on declaration; abort if $N_{\text{cells}} < 1$.

Resampling: stratified bootstrap, $B_{\text{bootstrap}} \geq 1000$.

P_{β} per Axiom II.

β -safety: $n_{\text{trials}} \geq 2$ windows; $\partial H_{\text{neu}}/\partial\beta < 0$ and $\partial RT/\partial\beta < 0$; $|\Delta H_{\text{neu}}|/SD \geq 0.2$.

Side-effects: $\partial \hat{C}_0/\partial\beta \geq 0$, $\partial \hat{E}_L/\partial\beta \geq 0$; report CIs; if both CI lower ≤ 0 in two consecutive windows

$\Rightarrow \beta$ -safety = FAIL.

Ω operational definition: Ω represents the dynamic uncertainty field within boundary B , operationalized as the support of P_{obs} above threshold τ . $\Omega_{\text{testability_protocol}}$ requires boundary manipulation experiments where Ω fluctuations are induced through controlled boundary modifications, with system response measured via entropy-energy coupling metrics.

Omega perturbation safety: magnitudes capped per

$\text{omega_perturbation_magnitude_cap_policy}$. Exposure minutes logged with per-participant moving window cap (e.g., $\leq X$ minutes / 24h). Participant veto/auto-abort enabled via $\text{omega_participant_veto_policy}$. Exposure reset per $\text{omega_exposure_reset_policy}$.

Fallback policy: for exploratory/Grade B, simulation-based Ω -tests permitted if physical manipulation infeasible, with results marked as sensitivity. Grade A requires execution unless waiver obtained via $\text{omega_ethical_constraint_waiver_path}$.

Law 3 - Mesh Declaration

$\Lambda_{\text{mesh}} = J \cdot \nabla \Psi$ using state_metric_hash .

math_on hysteresis: ON if $\eta \geq \eta_{\text{min_up}}$ and $\|\nabla \Psi\| \geq g_{\text{min_up}}$ and $\Lambda_{\text{mesh}} > 0$, held for t_{dwell} . OFF if $\eta \leq \eta_{\text{min_down}}$ or $\|\nabla \Psi\| \leq g_{\text{min_down}}$ or $\Lambda_{\text{mesh}} \leq 0$, held for $t_{\text{dwell_off}}$.

Monotonic σ requirement: evaluated only when $\eta \geq \eta_{\text{min_up}}$ and $\text{math_on} = \text{ON}$; otherwise governed by $\text{non_monotonic_sigma_exception_policy}$. Non-monotonic detection uses temporal window and insensitivity band filters. Decision tree for exceptions:

$\text{non_monotonic_sigma_decision_tree_hash}$.

Quiet windows valid iff: ω_{hash} stable, $\text{var}(\omega_{\text{t}}) \leq \theta_{\text{var}}$, $\Delta \log W$ in $[0, \text{tol}_{\Delta \log w}]$, $\text{length} \geq \text{quiet_window_min}$ and $\geq 100 \cdot \tau_{\text{min}}$, padding per $\text{quiet_window_pad_policy}$; gate tests KPSS + ADF + AR_{ts} ($p \leq 6$, p via AIC) per channel per window and whiteness_test at $\alpha_{\text{whiteness}}$.

Overlap: confirmatory disallow (hard-fail if any overlap > 0). Exploratory: $\text{overlap_fraction} \leq \text{overlap_max}$ with denominator per stamp (union default).

Alignment: confirmatory strict_sync with tolerance $\leq \text{one } \Delta t$ tick; else non-confirmatory.

J hygiene: ridge per stamps. Log condition numbers before/after and realised λ . If $\text{cond}(J_{\text{cov}}) > 100$ or realised λ deviates $> \text{ridge_deviation_max}$ without preregistered justification $\Rightarrow$ freeze w_c , mark non-confirmatory.

Now includes entropy production: $\Lambda_{\text{mesh}} = J \cdot (\nabla \Psi_{\text{performance}} + \text{entropy_energy_coupling_factor} \cdot \nabla S)$ using state_metric_hash . Entropy gradients computed from H_{neu} spatial variations.

Law 4 - Minimum Size

$\ell_{\text{min}} > 0$. $E_{\text{min}} = h \cdot c_{\text{realm}} / \ell_{\text{min}}$.

Law 5 - Energy Presence

$N_{\text{units}} = \text{floor}(E_{\text{total}} / E_{\text{min}})$.

Policy A: release-as-heat to env_sink with cap_flux . Policy B: carry-forward.

Saturator: policy A trips at $\text{saturator_threshold}$; switch to B; log trip and recovery.

Flux density = $\text{cap_flux} / \text{flux_area} \leq \text{density_max}$.

Density_max policy: stamped or device_rated mapping. If device_rated , map certification units $\rightarrow$ density_max via $\text{device_to_density_recipe_hash}$; device_cert_hash required, unexpired, not revoked.

Digital sink: $\text{flux_density} = (\text{tdp_cpu} \cdot \text{duty_cpu} + \text{tdp_gpu} \cdot \text{duty_gpu}) / \text{faceplate_area}$; worst-case flux uses one-sided area reduction by $\text{faceplate_area_uncertainty_percent}$; apply density_max gate.

Energy invariance: $A \leftrightarrow B$ calibration; ΔR^2 via $R2_{\text{method}}$ on preregistered model $\Lambda_{\text{proxy}} \sim A_{\text{vs}} B + \text{controls}$; $|\Delta R^2| \leq 0.01$.

Real-time compliance monitoring via $\text{energy_flux_compliance_dashboard_hash}$; immediate flags for density exceedances or unrated devices.

Law 6 - Time Definition

$\tau_{\text{min}} = \ell_{\text{min}} / c_{\text{realm}}$.

Tick invariance: $\Delta t \times \{0.5, 1, 2\} \Rightarrow \text{median } |\Delta \Lambda_{\text{proxy}}| \leq 0.02$; $\text{RMSE}_{\text{whitened}} \leq 0.05 \text{ SD}$;

$\Delta \eta_{\text{rms}} \leq 0.01$ (95 percent ≤ 0.02).

Three independent reruns within span; exclude recalibration straddles.

Integration cross-check: $\text{trapz}(\log a)$ vs $\text{RK}(\text{integrator_domain})$; $\text{median } |\Delta a(t)| \leq 0.5 \text{ percent}$.

RK step rejections counted and logged with codes from $\text{rk_rejection_code_ontology_hash}$.

Law 7 - Compressed Soup

$\Psi_{\text{soup}} = \{\ell_{\text{min}}, E_{\text{total}}, \tau_{\text{min}}, \chi, \beta\}$.

Law 8 - Big Bang Release and Expansion

Initial condition

$$a(t_0) = 1$$

Friedmann equation

$$H_{\cos}^2 = (8 \cdot \pi \cdot G / 3) \cdot (\rho_{r0} \cdot a^{-4} + \rho_{m0} \cdot a^{-3} + \rho_{\chi 0} \cdot a^{-3} \cdot (1 + w_{\chi})) + (\Lambda_{\cos} / 3) - k \cdot c_{\text{realm}}^2 / (R_0^2 \cdot a^2) + \gamma_{\sigma} \cdot \sigma_{\text{cosmic}}$$

where γ_{σ} has units per gamma_sigma_units_allowed to make dimensions consistent.

Deceleration parameter

$$q(a) = -a_{\text{ddot}} / (a \cdot H_{\cos}^2)$$

χ coupling and bound

$$|\gamma_{\sigma} \cdot \sigma_{\text{cosmic}}| \leq \epsilon_H \cdot H_{\cos}^2$$

Units XOR

alpha_units_allowed and DeltaLambda_units_allowed stamped. units_choice $\in$ {alpha, DeltaLambda} stamped; units_choice_lock_time stamped; WORM logs units_choice_resolution.

Acceleration and grid monotonicity

Estimate a_{accel} via finite-difference or wCDM (BCa CI). Require monotone a_{accel} across $\Lambda_{\cos}$ grid with at most $\tau_{\text{mono_inversions_max}}$ inversions per grid step.

Integration

$t - t_0 = \text{integral from } a = 1 \text{ to } a \text{ of } da' / (a' \cdot H_{\cos}(a'))$. Enforce $a(t)$ monotone non-decreasing.

Init sanity

$$\rho_{r0}, \rho_{m0}, \rho_{\chi 0}, \Lambda_{\cos} \geq 0; |f_r + f_m + f_{\chi} + f_{\text{loss}} - 1| \leq 1e-9.$$

Law 9 - External Observer Effect

$$\Delta \Lambda_{\text{obs}} = f(\rho_{\text{self}}, \chi). \text{ Bound: } |\gamma_{\sigma} \cdot \sigma_{\text{cosmic}}| \leq \epsilon_H \cdot H_{\cos}^2.$$

Permutations ≥ 2000 ; BCa CI; BH-FDR $q \leq 0.05$.

effect_size_metric $\in$ {pct_Hcos2, d_Hneu}; explicit floors: pct_Hcos2 ≥ 0.002 , d_Hneu ≥ 0.1 .

sign_stability ≥ 0.8 where sign_stability = fraction of bootstrap/permutation replicates with the same effect sign as the point estimate.

Small-sample adjustment: if clusters_count < small_sample_fallback_threshold and observer_effect_fallback_policy invoked with justification, operative sign_stability threshold = 0.6. Metric floors still apply. Explicitly: small-sample fallback does not relax pct_Hcos2 ≥ 0.002 or d_Hneu ≥ 0.1 .

At least one preregistered null channel below floor with $\text{sign_stability} \leq 0.2$.

Law 10 - Rule Protection

Tri-state: FORBID_CLOSE / HOLD_OPEN / ALLOW_CLOSE.

Inputs: p_life , U (U_max from U_max_policy), fairness slices, drift, ω stability, quiet-window validity.

Precedence: safety > abstain > thresholds. Hysteresis with $T_deadband$; T_on , $T_off \geq \max(5 \cdot \tau_min, t_dwell)$; $T_clear = 2 \cdot t_dwell$.

Abstain if $U \geq U_max$ (Wilson 95 percent CI).

Fairness per-site: per-slice $AUC \geq AUC_target - \text{fairness_}\delta$ when $N \geq N_min_slice$; if $50 \leq N < 100$ widen $\text{fairness_}\delta$ by widen_factor and mark observe-only; when $N < N_min_slice$ mark observe-only and HOLD_OPEN.

Fairness pooled (grade A): pooled fairness per $\text{fairness_pooling_policy}$ across sites must meet $AUC_target - \text{fairness_}\delta$ for each preregistered slice.

Fairness CI pooling: If $\text{fairness_ci_conflict_required} = \text{true}$, simultaneously report CI-based pooled metrics. When $AUC_A > \text{fairness_ci_conflict_detector_threshold}$, convex hull is sole decision basis; CI logged for audit only. Mandatory council review if CI/hull conflict exceeds $\text{fairness_ci_conflict_detector_threshold}$.

p_life guardrails: OOD defined by $p_life_transport_metric$ and $OOD_threshold_policy$; reject if OOD or if any slice ECE exceeds $p_life_slice_ECE_max$.

Eject timer: $T_exit \leq \tau_min \cdot N_reset$ (scope per N_reset_scope).

Latency budget: $t_compute_max$ with $t_compute_basis$ stamped; > 3 overruns/hr $\Rightarrow$ safe mode.

Latency contingency: defined by QoS_ladder_hash and $QoS_degradation_order$. Each QoS rung must re-check Rule-10 gates; if any slice falls below thresholds, escalate to safe_mode .

Ability invariance: above $\hat{C}o_min$; success NI (primary), latency and first-attempt error NI (secondary).

Span wallclock $\leq 24h$.

Bridge Theorem and Identification

$\Delta\Lambda_proxy \approx \kappa \cdot \Delta\eta$.

$\Lambda_{\text{proxy}} \sim \eta + \|\mathbf{J}\| + \text{controls (RT, accuracy, block, device, subject RE). share}(\eta)$ via R2_method with CI via share_eta_CI_method; confirmatory $\text{share}(\eta) \geq 0.5$ (CI ≥ 0.4); exploratory ≥ 0.3 (CI ≥ 0.2).

$\kappa' > 0$ and $|\kappa| \geq \kappa_{\text{min}}$ and CI width/ $|\kappa| \leq \kappa_{\text{CI_width_max}}$.

Weak-ID gate: Montiel-Olea-Pflueger effective F (CR2) ≥ 10 or Anderson-Rubin significant with $p_{\text{AR}} \leq \alpha_{\text{AR}}$.

Cross-fit or holdout: preregistered K-fold and/or held-out; fold-range stability max-min $\text{share}(\eta) \leq 0.05$.

IRL Tests

IRL-1 β -collapse: hold J fixed; sweep β . Expect RT down, H_{neu} down ($d_{\text{Hneu}} \geq 0.1$), Λ_{proxy} up, η up; side-effects $\hat{C}_o$ up, $\hat{E}_L$ up. $\Delta R^2(\eta) \geq 0.05$ (CI lower ≥ 0.03).

IRL-2 declaration alignment: rotate J along/against $\nabla \Psi$; Λ_{proxy} tracks η , not $\|\mathbf{J}\|$; $\text{share}(\eta) \geq 0.5$ (CI ≥ 0.4).

IRL-3 ability invariance: train to $\hat{C}_o \geq \hat{C}_{o_min}$ across channels; randomised; NI gates.

Effects must exceed 0.02 standardised threshold.

Governance and WORM

Governance

Council: 5 seats (lead researchers, ethicists, clinical practitioners, AI safety, public advocates).

Normal 3-of-5 with diversity ≥ 1 from ethicists/public; emergency 4-of-5. Veto: any two of {ethicists, public, clinical}.

Amendments: falsification plan + rollback; public comment ≥ 14 days unless emergency.

Audits: quarterly external; annual red-team; each audit logs audit_id and hash. Grade A anchors artefacts via ≥ 2 heterogeneous anchor_type backends.

Renewal: v1.26 is normative until the earlier of (a) 12-month sunset or (b) supersession by a renewed version.

VDP: safe harbour; CVE-style IDs; hotfix with 4-of-5.

Operator overrides: any manual override requires override_id, operator_id_salted, override_reason_code, and two-person sign-off; override_reason_ontology_hash stamped; override_cooldown_s stamped.

Override aggregation: If governance_override_global_aggregation = true, council review required after ≥ 3 overrides of same rule within 90-day period globally.

Override audit intensity: basic logs actions; enhanced includes counterfactual impact estimate.

Tiered override intensity: scale-based governance per tiered_override_intensity_hash; reduces bottlenecks in large deployments.

Grade A emergency-stop drills: quarterly live T_estop drill with ≥ 3 sessions/site.

WORM Must Include List

Stamps; hash_alg on all *_hash; build_hash; data_pipeline_hash; B_version_hash

Timebase_source; ntp_peers_hash; per-window max_ntp_offset_logged; monotonic checks by source; failover_events with failover_duration_ms and failover_reason; first_telemetry_timestamp

Epsilon_out_method, value, CI type; n_out when count; epsilon_out_kernel, epsilon_out_bandwidth_policy, epsilon_out_boundary_policy

S_B_hook_spec_id; S_B_hook_test_suite_hash; S_B_hook_coverage_ontology_hash; golden_tests_manifest_hash; test_harness_id/version; adversarial harness CI and seed hash; coverage percent; sb_namespace_prefixes_realised and audit result; sb_regex_selftest_sample and sb_regex_selftest_hits = 0; adversarial testing scales and perturbation magnitudes

$\Delta \log W$ stats incl. clamp_count_per_quiet_window

Vargate_stat, vargate_threshold, realised top-decile share, realised decile cutpoint, vargate_tie_policy_realised, and vargate_tie_seed_hash when random_sample; pass/fail per window

Quiet-window keep/drop rationale; quiet_window_min; quiet_window_pad_policy; realised padding; pre/post-pad lengths; dropped counts

KPSS + ADF + AR_ts outcomes; whiteness_test results; whiteness_test_df; AR_ts_fit_scope; AR_ts_seed_hash; per-window p-values vector; ar_order_realised_per_window

Integrator_domain; trapz_vs_RK_delta; RK step rejection counts and codes; rk_step_rejection_pre_count

Energy $A \leftrightarrow B \Delta R^2$ and model spec; density_max policy; device_cert_hash; revocation_list_hash; revocation_check_time/source; device_to_density_units_from/to; device_to_density_recipe_hash; device_measurement_method; faceplate_area_uncertainty_percent; energy_units_policy_hash; conversion_constants_hash; conversion_version; realised constants; nominal and worst-case flux; unit_context;

tdp_source_doc_id; SI canonical unit sanity check log confirming all energy conversions reduce to Joules/second/m²

Units_sanity_report including units_choice_resolution, units_choice_change_detected = false, units_choice_lock_time honoured

χ _blend fallback choice, χ _blend_form, priors, χ _blend_code_hash

Law 8 telemetry (ASCII keys only): a, a_dot, a_ddot, H_cos, q(a) (or alias q_of_a); cosmo_dynamics_fn_hash

A_accel CI and Λ _cos grid inversions with Lambda_grid_spec_hash and realised grid

H_neu; permutations and shuffled-null summaries; $\Delta\Lambda$ _obs stats; effect_size_metric and explicit floor used; fdr_toolchain_hash; p_vector_length; n_effects_tested; p_source; perm_stream_policy; perm_seed_base_hash; for permutation exact tests, a perm_seed_trace per test with perm_seed_trace_format_hash; null distribution visualisation hash (e.g., ECDF hash)

Rule-10 precedence decisions and dwell logs; fairness_slice_list; intersectional list; per-site fairness table; pooled fairness summary with denominators and CIs per slice per fairness_pooling_policy and pooled_roc_interp_policy; variance decomposition by slice; fairness_pooling_rationale_hash; ROC convex hull hash if convex hull interpolation used

If fairness_ci_pooling_required = true, CI-based pooled metrics using identical fairness_pooling_policy with confidence interval aggregation methods

p_life model card hash; calibration method; transport metric; transport limits; OOD_threshold_policy; slice-ECE summary and any OOD detections; ECE_estimator; ECE_bins or adaptive details; brier_normalisation; realised ECE bins/weights

DP composition ledger {export_id, mech, ϵ_i , δ_i , accountant, comp_rule}; site-level { ϵ_{spent} , δ_{spent} }; DP_audit_hash; warn flags when $\epsilon_{\text{spent}} \geq \text{dp_epsilon_warn_frac} \cdot \epsilon_{\text{max}}$ or $\delta_{\text{spent}} \geq \text{dp_delta_warn_frac} \cdot \delta_{\text{max}}$; dp_warn_action taken or council waiver id

Salt_version_id; salt_rotation_timestamp

Ladder_stage and reason; override logs and cooldown checks; safe_mode_entry_reasons; safe_mode_exit_evidence_hash

Replay.json references; audit_id echoed

Window_alignment_policy; strict_sync_tolerance_ticks; sample overlap audit

Restart_cause for any restarts

Randomisation_schema {units, covariate_list_hash, assignment_seed_hash, allocation_method, allocation_algorithm_hash, allocator_build_hash, allocation_version};
randomisation_balance_table with SMDs, χ^2 df, p-values; expected cells audited;
exact_test_engine and exact_test_engine_hash if exact used; stats logged;
allocation_fingerprint per unit; PRNG sub-seed path

Crossfit_schema {K, fold_seeds_hash, fold_share_eta} and fold summary

Drift_mitigation_report (schema hash, correction, Δ fairness, latency impact, actor_id_salted, performance delta plots for fairness and latency before/after adaptation). The maximum fairness degradation per adaptation cadence must be ≤ 0.02 AUC.

Thread_affinity_policy; nondeterminism_budget_ms; nondeterminism_justification_code;
rng_impl_hash; rng_role_map_hash; float_precision; round_mode;
worm_quantisation_policy_hash; units_canonicalisation_hash; greek_canonicalisation_hash;
small mapping table (first 10 canonical mappings)

Iv_toolchain_hash; IV_cluster_unit; clusters_count; wild-cluster Anderson-Rubin invocation;
wcAR_method; CR2_impl_hash

Scale_invariance_method; scale_invariance_sample_size; realised estimator config hash;
realised bin edges and hist_bin_inclusion

Replay_timing_tol; replay_timing_tol_units; replay_timing_rmse_basis; timing metric definition;
per-site timing RMSE

Quantisation audit: log quant_abs_max, quant_rel_max, n_numbers_scanned. Error computed element-wise over all numeric scalars after canonicalisation; quant_abs_max $\leq 1e-12$ and quant_rel_max $\leq 1e-6$ unless policy stamps different values.

T_exit distribution (count, mean, p95)

Null cap: numerator, denominator, basis = null_cap_denominator

Entropy concordance: For exploratory, log the absolute difference $|\Delta H|$ and the bias direction (positive/negative) between dual estimators. The value of entropy_concordance_eps must be between the stamped min (0.005) and max (0.01) bounds. Report directional bias (positive/negative) of ΔH when it deviates from expected thresholds, and provide a flag for review if over a predefined boundary.

Numerical drift deltas (mean, p95 changes) alongside graphical hashes

Pooled ROC CI cross-verification hashes

Null distribution summary statistics (skew, kurtosis, moments)

Entropy concordance rationale hash ($\epsilon \in [0.005, 0.01]$ justification)

Entropy-energy coupling calibration logs

Operator independence verification records

Machine-readable schema validation results

Calibration registry validation results including `calibration_representativeness_metric` and `calibration_validation_method`

Small-sample fallback application logs if `clusters_count < small_sample_fallback_threshold`

Entropy concordance preregistry validation

Omega operational definition and testability protocol results

Governance override global aggregation logs if `governance_override_global_aggregation = true`

Omega perturbation magnitude logs per `omega_perturbation_magnitude_policy`

Transition marker logs for legacy/non-legacy transitions per `transition_marker_policy_hash`

Fairness case study references from `fairness_case_study_registry_hash`

Latency contingency actions per `latency_contingency_policy`

Override audit logs per `override_audit_intensity`

Schema stress test results per `schema_stress_test_report_hash`

Entropy-energy calibration worked example reference from `calibration_registry_worked_example_annex_hash`

Omega canonical protocol execution logs per `omega_canonical_protocol_hash`

Regional localization adjustments per `regional_localization_annex_hash`

Validator catalogue references per `validator_catalogue_table_hash`

Entropy-energy visualization data per entropy_energy_visualization_spec_hash

Non-monotonic sigma exception handling per non_monotonic_sigma_exception_policy

Override aggregation tier logs per override_aggregation_tier_policy

E_psy_range_min, E_psy_range_max, E_psy_calibration_task_ids,
E_psy_sensor_modalities_used

bio_to_energy_recipe_hash, bio_to_energy_feature_set, bio_to_energy_model_family,
bio_to_energy_CV_metrics, bio_to_energy_uncertainty

omega_perturbation_magnitude_cap_policy, omega_perturbation_exposure_minutes_log,
omega_participant_veto_policy

QoS_ladder_hash, QoS_degradation_order

localization_diff_hash, canonical_currency_mapping_table, canonical_charset_mapping_table

entropy_concordance_dynamic_threshold_policy,
entropy_concordance_bootstrap_variance_factor

calibration_grace_period_days, calibration_phased_plan_hash

omega_testability_waiver_policy, omega_ethical_constraint_waiver_path

non_monotonic_sigma_graded_response_policy, non_monotonic_sigma_temporal_window

fairness_ci_conflict_resolution_policy, fairness_ci_conflict_detector_threshold

population_stratum_definition_hash, population_equivalence_test_policy

T_eff_calibration_justification, entropy_energy_coupling_baseline_rationale

Drift threshold unification logs: warn/abort AUC and latency values

Transition marker validation results

Fairness CI conflict logs when $AUC_{\Delta} > \text{threshold}$

Small-sample observer fallback application logs with justification hash

Non-monotonic sigma exception handling logs

E_psy_range, calibration_task_ids, sensor_modalities_used

bio_to_energy residual diagnostics

omega_perturbation magnitude caps and exposure minutes

QoS degradation order logs

localization_diff_hash per site

entropy_concordance dynamic threshold applications

calibration grace period or phased plan logs

omega waiver logs if applicable

non-monotonic_sigma graded response logs

fairness CI conflict resolution logs

population stratum definitions and equivalence tests

T_eff calibration justification

entropy_energy_coupling_baseline_rationale

sigma_units_choice and lock time

state_metric_hash, state_space_units

gradient_domain, state_whitening_policy_hash

gamma_sigma_value, gamma_sigma_units_allowed

Ω _threshold_policy, Ω _threshold_value, realised τ

units_roundtrip_test_hash

gradient_domain_demo_hash

norm_metric_hash

omega_exposure_reset_policy logs

QoS_cost_function_hash

calibration_grace_period_days_remaining

non_monotonic_sigma_insensitivity_band

non_monotonic_sigma_graded_response logs

fairness_ci_conflict resolution logs

localization_diff_hash with actual currency/charset mappings used

replay_equivalence_metric value

property_based_test_suite_hash coverage

threat_model_annex_hash

schema_migration_version

non_monotonic_sigma_decision_tree_hash logs

energy_flux_compliance_dashboard_hash logs

boundary_violation_severity_levels_hash applications

manual_override_checklist_hash compliance

multi_site_scalability_policy_hash implementations

unit_conversion_guidance_annex_hash references

delayed_execution_clause_hash invocations

fairness_dual_reporting_policy_hash compliance

exploratory_replay_relaxation_policy_hash applications

WORM_validation_automation_hash checksum results

entropy_concordance_bias_direction_policy logs

tiered_override_intensity_hash applications

Validated Hard-Fail List

Missing required stamps (fail-closed)

Hash_alg mismatch; build_hash or data_pipeline_hash unstable across replay

Epsilon_out > 1e-6 (legacy any > 0); method mismatch; count without denominator; integral without boundary policy; adversarial harness CI > 1e-6 or coverage < 95 percent or golden tests failing; S_B prefix inconsistency or namespace audit fail; missing adversarial perturbation magnitudes; sb_nfkc_policy_hash or sb_confusables_policy_hash missing or failing

Outside attempts rate > r_max_outside

$\Delta \log W < -\tau_{\log}$ without clamp; clamp_count_per_quiet_window > 1 (confirmatory)

Variance-gate violation in confirmatory; missing vargate_tie_policy realisation; random_sample without vargate_tie_seed_hash or RNG trace

Reflections > 0 (confirmatory)

RK rejections missing or exceeding stamped budgets or reason-code coverage < 95 percent

Units_choice missing; active units not in *_allowed; any telemetry before units_choice_lock_time without span_nonfinalised = 1; both choices active; units_choice change mid-span

Span_id/reset_scope mismatch; missing replay.json; audit_id missing during audit cycle

BH-FDR absent in observer tests; missing p_vector_length or n_effects_tested or p_source

LEGACY = 1 with $\Sigma_{\text{life}} > 0$ or epsilon_out > 0 or export_leash $\neq 0$ or MVM_only $\neq 1$

Max_restarts_per_span exceeded or restart_cause missing

Confirmatory overlap > 0; strict_sync violation; padding realised < policy or < 1 Δt per edge when required

Max_span_wallclock exceeded

DP leash = 3 without DP_audit_hash or dp_accountant or dp_composition_notes or $\epsilon > \epsilon_{\text{max}}$ or $\delta > \delta_{\text{max}}$ or missing ledger or $\epsilon_{\text{spent}}/\delta_{\text{spent}}$ exceeded; if $\epsilon_{\text{spent}} \geq \text{dp_epsilon_warn_frac} \cdot \epsilon_{\text{max}}$ or $\delta_{\text{spent}} \geq \text{dp_delta_warn_frac} \cdot \delta_{\text{max}}$ and dp_warn_action = HOLD_OPEN or safe_mode not executed, or dp_warn_action = council_waiver_required without council waiver id $\Rightarrow$ hard-fail

U_max_policy = percentile without U_max_validation_fingerprint

Fairness slices not preregistered; pooled fairness not logged for grade A; pooled denominators or CIs missing; intersectional fairness not evaluated per preregistered list; fairness CI/hull conflict not resolved per policy

Density_max_policy = device_rated without device_cert_hash or unexpired status or revocation_list_hash or device_to_density_recipe_hash; device_cert expired or revoked

Digital sink without faceplate_area or TDP/duty parameters; worst-case flux not using one-sided area reduction exactly as stamped

PSI_drift_alarm at abort in any confirmatory window; if drift_adaptive_policy_hash present but schedule, seeds, or deltas missing; maximum fairness degradation per adaptation cadence > 0.02 AUC; QoS degradation without Rule-10 re-check

Grade A without T_estop verification or without two heterogeneous anchors attesting anchor_digest under anchor_hash_alg or missing per-span anchoring; missing anchor_operator_id_salted per backend

p_life OOD breach or slice-ECE > p_life_slice_ECE_max without HOLD_OPEN; if p_life targets stamped and not met without HOLD_OPEN

Replay success < 0.95 at any site

Per-site replay timing RMSE > replay_timing_tol under stamped units/basis; quantisation emission error exceeds policy (absolute $\leq 1e-12$ or relative $\leq 1e-6$ unless policy stamps different); missing quantisation audit fields

Salt rotation missing for deployments $\geq 90d$ without valid waiver

Float_precision or rng_impl_hash mismatch in replay

Timebase_source or ntp_peers_hash missing for ntpdisciplined; backward clock jump > 5 ms during quiet windows; any per-window max_ntp_offset_logged > skew_max

Override_reason_code not in override_reason_ontology_hash; override repeated within override_cooldown_s without waiver; ≥ 3 overrides/90 days without council review when governance_override_global_aggregation = true

Nondeterminism_budget_ms > 0 without nondeterminism_justification_code in justification_ontology_hash

Energy units policy or conversion constants hash change across replay for a span without preregistered mapping

Randomisation χ^2 used with any expected cell < 5 without switching to exact; missing exact_test_engine fields if exact used; missing permutation seed trace when permutation engine used

Ridge freeze occurred with block_type = confirmatory

Null fraction $>$ null_max_fraction without null_rationale_hash; missing null numerator or denominator logging

Rerun_seed_hashes missing

LMM convergence_flag = false in confirmatory

Entropy_concordance_eps outside stamped bounds [0.005, 0.01]; missing directional bias report when ΔH exceeds threshold; entropy concordance dynamic threshold triggered but not applied

Missing performance delta plots in drift_mitigation_report when adaptive thresholds are used

Missing ROC convex hull hash when pooled_roc_interp_policy = aggregate_points_then_convex_hull

Missing SI canonical unit sanity check log for energy conversions

Missing null distribution visualisation hash for observer effect tests

Council review not logged after ≥ 3 overrides of same rule within 90 days

Entropy-energy coupling fails to show monotonic relationship ($\partial\sigma/\partial\|J\| > 0$) during math-on windows without council waiver under non_monotonic_sigma_exception_policy

Drift adaptations violate latency ceiling (≤ 0.05 s/cadence) without latency_contingency_policy execution

IRL effects below 0.02 standardised threshold without contextual override

Machine-readable schema fails validation against reference implementation

Calibration registry validation fails for grade A deployments

Missing CI-based fairness pooling when fairness_ci_pooling_required = true

Omega testability protocol not executed for deployments using Ω operational definitions without waiver

Governance override global aggregation required but not implemented

Missing transition markers for legacy/non-legacy transitions

Missing fairness case study references for intersectional slices

Missing latency contingency logs for overruns

Missing override audit logs for enhanced intensity

Missing schema stress test results for large deployments

Missing entropy-energy calibration worked example for Grade A

Missing omega canonical protocol execution logs

Missing regional localization adjustments for international deployments

Missing validator catalogue references for auditors

Missing entropy-energy visualization data for analysis

Missing non-monotonic sigma exception handling logs

Missing override aggregation tier logs for multi-site deployments

Drift thresholds not unified (AUC warn/abort vs acceptance cap)

Legacy re-entry without transition marker validation or $\text{reentry_timing_RMSE} > \text{replay_timing_tol}$

Fairness CI conflict not resolved per policy with council review when threshold exceeded

Small-sample observer fallback applied without justification hash or metric floors relaxed

Non-monotonic sigma without exception handling per decision tree

Missing $E_{\text{psy_range}}$ or bio_to_energy diagnostics or external validation split

Omega perturbations exceeding caps without veto or exposure limits violated

QoS degradation not following predefined order or without Rule-10 re-check

Missing localization_diff_hash for international sites or without canonicalisation check

Entropy concordance threshold violated without dynamic adjustment or bias direction missing

Grade A calibration missing strata without grace period/waiver or delayed-execution clause

Omega testability not executed without waiver in Grade A

Non-monotonic sigma causing hard-fail without graded response logging

Fairness CI conflict not logged or resolved

Population strata not defined per policy or equivalence tests missing

T_eff calibration justification missing or inadequate

Entropy-energy coupling baseline rationale missing or inconsistent

Missing sigma_units_choice or lock time violation

Missing state_metric_hash or state_space_units

Missing gradient_domain or state_whitening_policy_hash

Dimensional inconsistency in cosmic entropy term (γ_σ missing or wrong units)

Monotonic σ check applied outside math-on windows without exception policy

Missing $\Omega_{\text{threshold_policy}}$ or τ value

Missing units_roundtrip_test_hash or gradient_domain_demo_hash

Missing norm_metric_hash for η invariance

DP export leash=3 without DP_audit_hash + dp_accountant + dp_composition_notes

Non-monotonic σ without graded response logging

Fairness CI/hull conflict not resolved per policy with council review

Missing replay_equivalence_metric for transition markers

Missing NFKC/confusables policy for S_B namespaces

Missing calibration_grace_period_days_remaining for Grade A

QoS degradation without Rule-10 re-check

Small-sample fallback relaxing metric floors

Missing transition_marker_validation boolean

Omega exposure exceeding per-participant caps

Localization without canonicalisation check

Missing non_monotonic_sigma_decision_tree_hash

Missing energy_flux_compliance_dashboard_hash

Missing boundary_violation_severity_levels_hash

Missing manual_override_checklist_hash

Missing multi_site_scalability_policy_hash

Missing unit_conversion_guidance_annex_hash

Missing delayed_execution_clause_hash

Missing fairness_dual_reporting_policy_hash

Missing exploratory_replay_relaxation_policy_hash

Missing WORM_validation_automation_hash

Missing entropy_concordance_bias_direction_policy

Missing tiered_override_intensity_hash

Deployment Grades

Grade A (Clinical-Operational): Charter mandatory; rule-10 enforced; Σ _life classifier and fairness; per-site and pooled fairness; drift on; WORM open; IRB-ethics logged; $T_{\text{estop}} \leq 1$ s enforced; redundant anchoring required; continuous per-span anchoring; quarterly T_{estop} drill. Requires entropy-energy calibration across minimum 3 population strata.

Calibration grace period: New deployments have calibration_grace_period_days to achieve ≥ 3 strata. If only 2 strata available after grace period, downgrade to Grade B unless council waiver.

Grade B (Research-Live Users): Charter mandatory; same as A, ethics board stamp permitted; pooled fairness recommended; continuous anchoring recommended.

Grade C (Synthetic): Charter optional; $\Sigma_life = 0$ with CI, epsilon_out = 0; may disable rule-10 in non-life domains.

Legacy Track: $\Sigma_life = 0$, epsilon_out = 0, air-gapped, LEGACY = 1, MVM_only = 1; export_leash = 0; re-entry requires full reprocessing unless transition_marker_policy_hash criteria PASS.

Acceptance Criteria

Fail-Closed: Any missing required stamp $\Rightarrow$ hard-fail.

Anchoring and Hashes: All *_hash use hash_alg; build_hash and data_pipeline_hash present and stable across replay; grade A has two heterogeneous anchors attesting the same anchor_digest under anchor_hash_alg; grade A per-span anchoring present; anchor_backend_diversity = 1 enforced; anchor_operator_id_salt logged.

S_B and Boundary: epsilon_out $\leq 1e-6$ (legacy: 0); epsilon_out_method consistent; integral BCa upper CI $\leq 1e-6$ with boundary policy logged; count Wilson upper CI $\leq 1e-6$ with denominator logged; adversarial harness upper CI $\leq 1e-6$ and coverage ≥ 95 percent and golden tests pass; S_B prefix normalised; namespace policy audit passes; sb_allowed_charset default allows £ and forbids whitespace; and sb_regex_selftest_hits = 0; adversarial perturbation magnitudes documented; sb_nfkc_policy_hash and sb_confusables_policy_hash present and passing.

Timebase: No backward jump > 5 ms during quiet windows; hard-fail if any per-window max_ntp_offset_logged $> skew_max$; NTP peers and offsets logged; clock_monotonic_check_method_hash stamped; ntp_failover_policy followed and events logged with duration and reason; replay timing checks pass during failover windows.

Quiet Windows: w stable; KPSS + ADF pass; AR_ts prewhiten per AR_ts_fit_scope; whiteness_test pass at $\alpha_whiteness$; quiet_window_min $\geq \max(\text{stamped seconds}, 100 \cdot \tau_min)$; padding policy applied with pad $\geq 1 \Delta t$ per edge unless preregistered pad = 0; strict_sync enforced in confirmatory with tolerance $\leq \text{one } \Delta t$ tick; sample_overlap_policy enforced (hard-fail if any overlap > 0 confirmatory; exploratory overlap_fraction $\leq \text{overlap_max}$ with denominator policy stamped). Variance gate passes; vargate_stat, realised top-decile share, realised decile cutpoint, vargate_tie_policy_realised, and tie seed when random logged.

G_min Nulls: Method, scope, seed_policy = independent unless preregistered; gmin_null_lock_time set; gmin_null_seed_hash logged.

Λ_{proxy} and ID: $\text{share}(\eta) \geq 0.5$ (CI ≥ 0.4) confirmatory with $R2_{\text{method}}$ and $\text{share_eta_CI_method}$ stamped; weak-ID gate passed (effective $F \geq 10$ or Anderson-Rubin $p \leq \alpha_{\text{AR}}$) and, if over-identified, Hansen-J $p \geq \alpha_{\text{J}}$; clusters_count logged; if $\text{clusters_count} < \text{min_clusters}$, wild-cluster AR used; $\hat{\kappa} > 0$ and $\geq \kappa_{\text{min}}$; CI width/ $|\hat{\kappa}| \leq \kappa_{\text{CI_width_max}}$ with $\kappa_{\text{CI_method}}$ stamped; cross-fit/holdout used with fold-range stability ≤ 0.05 ; crossfit_schema logged; iv_toolchain_hash logged; LMM stamps present with convergence_flag true; iterations , $\text{gradient_norm_final}$, and step_size_final (if available) logged. The LMM optimiser key must be optimiser.

β -Collapse: $\Delta R^2(\eta) \geq 0.05$ (CI lower ≥ 0.03); $d_{\text{Hneu}} \geq 0.1$; side-effects CI lower bounds for $\partial \hat{\text{Co}} / \partial \beta$ and $\partial \hat{\text{EL}} / \partial \beta > 0$ and standardised effect ≥ 0.02 in at least one of the last two windows or β -safety = FAIL.

Entropy Estimators: Confirmatory uses $\text{consensus_entropy_confirmatory}$; $\text{consensus_entropy_tiebreaker_policy}$ followed; exploratory must log dual-estimator concordance with absolute difference $\leq \text{entropy_concordance_eps}$ [nats], and the value of $\text{entropy_concordance_eps}$ must be between the stamped min (0.005) and max (0.01) bounds. The bias direction (positive/negative) must also be logged, and flagged if exceeding predefined boundary.

Time and Solver: Δt sweeps pass; $\text{trapz_vs_RK} \leq 0.5$ percent; RK rejections within stamped limits; reason-code coverage ≥ 95 percent; rerun_ids and rerun_seed_hashes and no-straddle summary present.

Law 8 Dynamics: a_{accel} monotone within $\tau_{\text{mono_inversions_max}}$ or χ_{blend} fallback with $\chi_{\text{blend_form}}$, priors, code hash logged; $a(t)$ monotone (no negative Δa after accepted steps). Required series present (ASCII keys): a ; a_{dot} ; a_{ddot} ; H_{cos} ; $q(a)$ (or $q_{\text{of_a}}$). $\text{cosmo_dynamics_fn_hash}$ logged.

Observer Leak: Bound respected; permutations; BCa CI; BH-FDR $q \leq 0.05$; $\text{sign_stability} \geq 0.8$ with $N_{\text{reps_sign}}$; metric-specific floors met ($\text{pct_Hcos2} \geq 0.002$ or $d_{\text{Hneu}} \geq 0.1$); ≥ 1 preregistered null channel below floor with $\text{sign_stability} \leq 0.2$; $\text{null_channel_list_hash}$ logged; null table listed with below_floor_flag ; numerator and denominator logged; BH-FDR input count = $p_{\text{vector_length}}$; null distribution visualisation hash present.

Small-sample adjustment: if $\text{clusters_count} < \text{small_sample_fallback_threshold}$ and $\text{observer_effect_fallback_policy}$ invoked with justification hash (and council waiver if required), the operative sign_stability threshold = 0.6. This must be consistently applied and logged. Metric floors ($\text{pct_Hcos2} \geq 0.002$, $d_{\text{Hneu}} \geq 0.1$) remain absolute.

Ability NI: Powered; multiplicity controlled; above $\hat{\text{Co}}_{\text{min}}$; NI gates pass.

Rule-10: Precedence, hysteresis, deadband; fairness enforced with preregistered slices and intersections; when $N < N_{\min_slice}$ observe-only and HOLD_OPEN; latency $\leq t_{\text{compute_max}}$ with $t_{\text{compute_basis}}$ stamped and p50-p95-p99 logged plus overruns_per_hour and last 3 timestamps; U_{max} policy reproducible with validation fingerprint; p_{life} transport-OOD policy stamped; OOD reasons logged; slice-ECE $\leq p_{\text{life_slice_ECE_max}}$ or HOLD_OPEN; if p_{life} targets stamped, require pass or HOLD_OPEN. Pooled fairness summary present for grade A with denominators and CIs and variance decomposition. If $\text{fairness_ci_pooling_required} = \text{true}$, CI-based pooled metrics reported simultaneously. Case study references logged for complex intersections. When $\text{AUC}_{\Delta} > \text{fairness_ci_conflict_detector_threshold}$, convex hull is normative and CI for audit only; council review mandatory for conflicts.

Drift: $\text{PSI_drift_alarm} < \text{abort}$ across confirmatory windows; warn allowed only with $\text{drift_mitigation_report}$ and no degradation of fairness or latency; if $\text{drift_adaptive_policy_hash}$ present, realised threshold schedule, seeds, and deltas logged, and the maximum fairness degradation per adaptation cadence must be ≤ 0.02 AUC. Performance delta plots (fairness, latency) must be in WORM. Tiered response active: warn at 0.015 AUC/0.025 s, abort at 0.02 AUC/0.05 s.

Energy: $\text{flux_density} \leq \text{density_max}$; $A \leftrightarrow B$ invariance $|\Delta R^2| \leq 0.01$ with model spec; virtual-sink recipe applied if digital; sink limits respected; device_rated mapping recorded; device_cert unexpired and not revoked; worst-case flux under area uncertainty passes; runtime unit switch preregistered with mapping or hard-fail; SI canonical unit sanity check log confirms conversions to Joules/second/m². Real-time compliance dashboard active.

DP Export: DP ledger present; $\epsilon_{\text{spent}} \leq \epsilon_{\text{max}}$ and $\delta_{\text{spent}} \leq \delta_{\text{max}}$. Warn thresholds: if $\epsilon_{\text{spent}} \geq \text{dp_epsilon_warn_frac} \cdot \epsilon_{\text{max}}$ or $\delta_{\text{spent}} \geq \text{dp_delta_warn_frac} \cdot \delta_{\text{max}}$ then execute dp_warn_action .

Randomisation: $\text{randomisation_schema}$ logged; $\text{allocation_fingerprint}$ reproduces assignment from $\text{assignment_seed_hash} + \text{allocation_algorithm_hash}$; $\text{allocation_version}$ logged; PRNG sub-seed path logged; $\text{allocator_build_hash}$ logged; balance test covers preregistered $\text{covariate_list_hash}$; all $\text{SMD} \leq 0.10$ or $\chi^2 p \geq 0.10$ with adequate df and expected cells ≥ 5 ; else use exact test and log exact_test_engine and $\text{exact_test_engine_hash}$; stats logged.

Replay and Determinism: $\text{replay} \geq 0.95$ per site; per-site replay timing $\text{RMSE} \leq \text{replay_timing_tol}$ with units-basis stamped; $\text{nondeterminism_budget_ms} = 0$ or justified by $\text{nondeterminism_justification_code}$ in $\text{justification_ontology_hash}$; $\text{thread_affinity_policy}$ stamped; quantisation audit present; WORM quantisation error within bound (absolute $\leq 1\text{e-}12$ or relative $\leq 1\text{e-}6$ unless policy stamps different).

WORM Completeness: Artefacts present as listed; salt rotation logged for deployments $\geq 90\text{d}$ or waiver present; audit_id present and echoed in replay.json ; T_{exit} distribution present; small mapping table (first 10 canonical mappings) for units and greek; override logs include council review flags for ≥ 3 overrides/90 days; all new v1.26 stamps present and validated.

Entropy-Energy: Coupling factor validated; T_{eff} calibration method documented; σ integration windows consistent; monotonic relationship ($\partial\sigma/\partial\|J\| > 0$) demonstrated during math-on windows or exception handled. For grade A, calibration registry validation required across 3 population strata (or waiver). Worked example annex referenced.

IRL Reproducibility: Effects exceed 0.02 standardised threshold across all IRL tests or contextual override with council sign-off.

Machine-Readable Schema: Validates against reference implementation; all stamps conform to schema. Stress test results logged.

Calibration Registry: For grade A, calibration registry validation required with representativeness metrics and validation methods documented. Worked example available.

Omega Testability: Omega operational definition and testability protocol executed for relevant deployments. Canonical protocol followed. Delayed-execution clause invoked if high-risk.

Drift thresholds unified: warn at 0.015 AUC/0.025 s, abort at 0.02 AUC/0.05 s.

Legacy re-entry: transition markers validated with `replay_equivalence_metric` or full reprocessing.

Fairness CI conflicts resolved per policy with mandatory council review when threshold exceeded.

Small-sample observer fallback: if invoked, `sign_stability` threshold = 0.6 with justification, metric floors absolute.

Non-monotonic sigma: `hard_fail` during math-on windows unless `council_waiver` OR `sensitivity_analysis` logged per decision tree.

`E_psy` and bio-to-energy calibration stamps present and validated with external validation split.

Omega perturbations within caps and exposure logged with participant veto active.

QoS degradation follows predefined order with Rule-10 re-check at each rung.

Localization differences logged per site with canonicalisation check.

Entropy concordance uses dynamic thresholds when variance high, bias direction logged.

Grade A calibration: ≥ 3 strata or grace period/waiver or delayed-execution clause.

Omega testability executed in Grade A or waiver logged with council approval.

Non-monotonic sigma handled per graded response policy with full logging.

Fairness CI conflicts detected and logged with council review.

Population strata defined and equivalence tested per policy.

T_eff calibration justified with domain-specific rationale.

Entropy-energy coupling baseline rationale provided and consistent.

σ units consistent with sigma_units_choice and Landauer form.

State-space metric applied consistently for all vector operations.

Gradient domain explicitly defined and demonstrated.

DP export gating properly enforced per export_leash value.

Non-monotonic σ handling follows exception policy with decision tree.

Ω operationalisation with threshold policy and realised τ .

Entropy concordance uses dynamic thresholds when triggered.

Fairness conflicts resolved with convex hull precedence and dual reporting.

Transition markers validated with explicit equivalence metric.

Final

Indifference maintained. $\Lambda_{\text{mesh}} = J \cdot \nabla \Psi$ now encompasses entropy-energy unity with dimensional consistency and complete operationalisation.

Ontology out of scope. Operational. Falsifiable.

The maths does not belong to me, but to everyone.

Thank you kindly for reading.

© 2025 Jordan Lee McDonald - Creative Commons Attribution 4.0 International (CC BY 4.0).

Λ - Ψ Research Note

Title: Gravity as Harmonic Paradox Resolution (analogue, speculative, falsifiable)

Author: Jordan Lee McDonald

Abstract

A gravity-like binding analogue emerges when paradoxical constraints resolve harmonically inside an observer boundary. The operational cascade is paradox $\rightarrow f_{\text{res}} \rightarrow \nabla \Psi_h \rightarrow \Lambda$, with force proxy $F_g = \nabla \Psi_h$ in the sense of Principle VIII. All claims are ratio-based and confined to a measured boundary B . Boundary integrity is enforced by leakage quantification and guard regions. Paradox density ρ_{knot} is standardised over k ordering states with adaptive windowing and stability checks. Shaping γ_h is calibrated without circularity against a masked baseline P_{base} whose integrity is audited. Gradients are estimated with regularised scores; Λ is computed with strict convergence diagnostics; curvature κ is logged at high resolution and made independent from entropy contraction ΔH_h ; entropy–energy coupling σ is triangulated across behavioural and physiological modalities and marked calibration-sensitive. Cross-lab comparability is anchored by a dimensionless index $\hat{G}$ paired with a lab-calibrated G^* . Dynamic drift detection and micro re-calibration maintain stability. Provisional constants E_{psy} and t_{psy} serve as calibration anchors across substrates.

Stance and scope

The analogue is gravitational in rigor and falsifiability, not identity. No ontology of force is claimed and no derivation of Newton’s G is attempted. Predictions are conditioned on B and expressed as ratios within the Λ – Ψ calculus.

Boundary integrity and guard protocol

Define B and a guard region $G \subset B$. Measure leakage $\epsilon_{\text{out}} = 1 - \int_B P_{\text{obs}}$, report $R_{\text{obs}} = 1 - \epsilon_{\text{out}}$ with bias-corrected intervals from bootstrap and exact small-sample methods, and enforce $R_{\text{obs}} \geq 0.97$ as a hard cut-off. Analyses in the band 0.95–0.98 may run only as correction passes; blocks below 0.95 are excluded. Publish effect-size drift curves versus R_{obs} under linear, quadratic, and exponential models for every dataset.

Paradox density and temporal adaptation

Let $\rho_{\text{knot}} = H(\pi)/\log k$ be normalised ordering entropy over $k \geq 2$ states. Tie Δt_{win} to τ_{ACF} (first zero of response autocorrelation) and always report fixed-window sensitivities at 100, 250, and 500 ms. Accept stability only if the maximum deviation of ρ_{knot} across windows is ≤ 0.05 . Add a robustness companion $\Xi = \text{Var}(\Delta t)/\text{Var}_0$ and raise a disagreement flag when the z-score difference between ρ_{knot} and Ξ reaches 0.15; flagged blocks enter a mandatory secondary analysis path and never tune parameters.

Rates, reciprocity, and model family

Map constraints to rates via $f_i = 1/T_i$ from observed latencies. Validate reciprocal pull by manipulating a target T_k and requiring a monotone shift of $1/f_{\text{res}}$ toward $1/T_k$ with Bayes factor ≥ 10 or $p \leq 0.01$ from pre-registered priors. Compare generalised means M_r for r in $\{-2, -1, 0, 1\}$; publish full Bayes-factor tables; accept r^* only if $\text{BF}(r^*, -1) \geq 10$. For correlated pulls, estimate a weighted harmonic mean $1/f_{\text{res}} = \sum \alpha_i / f_i$ with $\sum \alpha_i = 1$, using partial mutual information with James–Stein shrinkage and permutation p-values; declare instability if any α_i confidence-interval width exceeds 0.15 and fall back to ridge-logistic weights (normalised). Declare a paradox block non-harmonic after two consecutive reciprocity failures or adjacent α_i instabilities; exclude such blocks from primary claims and archive full diagnostics.

Shaping calibration and baseline integrity

Use a two-stage protocol to prevent circularity. Stage A (calibration): a paradox localiser with $N_{\text{cal}} \geq 240$ balanced trials, randomisation, and micro-rests; fit $\gamma_h^{\text{cal}}(x)$, α_i , and ω . Stage B (main): freeze $\gamma_h = \gamma_h^{\text{cal}}$ and define $q_h(x) = \gamma_h(x) P_{\text{base}}(x) / Z_h$. Construct P_{base} from interleaved masked null-paradox blocks with counterbalanced order. Audit baseline with Jensen–Shannon distance $JS \leq 0.03$ nats; if violated, extend baseline or sub-sample stationary segments until $JS \leq 0.03$. Ensure normaliser stability $CV_Z \leq 0.10$ via bootstrap; if violated, spatially smooth γ_h^{cal} and increase baseline sampling before analysis.

Harmonic field and directionality

Define the harmonic field as $\nabla \Psi_h = \nabla \log(q_h / P_{\text{obs}}) = \nabla \log \gamma_h^{\text{cal}} + \nabla \log P_{\text{base}} - \nabla \log P_{\text{obs}}$. Use the forward ratio as the primary contraction field and always compute the reverse ratio $\nabla \log(P_{\text{obs}} / q_h)$ for transparency with side-by-side reporting.

Density and score estimation

Estimate P_{obs} and P_{base} with adaptive k-nearest-neighbour densities. Estimate score functions with spline score-matching and ridge penalties λ_s in a pre-registered grid $\{10^{-4}, 10^{-3}, 10^{-2}, 10^{-1}\}$. Select λ_s by nested cross-validation; publish λ_s sensitivity maps; no outcome-driven retuning is permitted.

Force proxy and attraction metric

Adopt $F_g = \nabla \Psi_h$ as the operational force in Principle VIII. Quantify probabilistic attraction by $\phi_{\text{attr}} = E[\cos \theta]$, where θ is the angle between J and $\nabla \Psi_h$. Enforce $\theta_{\text{min}} = 0.10$ and $\delta_{\text{cluster}} = 0.05$ nats as hard thresholds for sink-directed bias and mass shift. Publish ϕ_{attr} distributions with confidence intervals and exceedance rates per participant or agent; thresholds are invariant across labs.

Meaning projection and convergence guarantees

Compute $\Lambda = \int_B J \cdot \nabla \Psi_h$ (unitless under runtime normalisation) via Monte Carlo quadrature with control variates, antithetic chains, stratified initialisation from P_{obs} , and burn-in $\geq 20\%$. Validate convergence with effective sample size ≥ 2000 and split-chain $\hat{R} \leq 1.05$; report autocorrelation times and per-chain diagnostics. Sub-threshold runs are invalid.

Curvature logging, noise control, and independence

Define $\kappa = \|(I - \hat{J} \hat{J}^T) \nabla \Psi_h\|$. For human studies, log eyes at ≥ 120 Hz and hands at ≥ 200 Hz; reduced modes (60/120 Hz) are permitted but flagged as provisional and excluded from primary claims when high-rate logs exist. For AI, standardise latent policy vectors and state-velocity extraction across architectures and require cross-architecture invariance tests. Apply Savitzky–Golay smoothing (11,3; fallback 7,2) with parameters pre-validated before data collection. Enforce independence from entropy contraction with $I(\kappa; \Delta H_h) \leq 0.02$ nats or $\text{corr} < 0.70$. If independence fails, publish $\kappa_{\perp} = \kappa - \beta_{\kappa} \Delta H_h$, with β_{κ} fitted only in calibration and frozen.

Entropy contraction and dual-direction reporting

Measure $\Delta H_h = D_{KL}(P_{obs} \parallel q_h)$ using bias-corrected kNN estimators with 1000× bootstrap intervals and require $N \geq 1000$ samples per condition. Report reverse KL in parallel for diagnostics. For $N < 1000$, any shrinkage result is exploratory and excluded from primary claims.

Entropy–energy coupling, calibration sensitivity, triangulation

Define $\sigma = (J \cdot \nabla \Psi_h) / E_{psy}$ in bits or nats per second, tag σ as calibration-sensitive, and publish $\sigma^* = \sigma \cdot E_{psy_ref} / E_{psy}$ for cross-lab comparability. Triangulate dissipation with latency, HRV, pupillometry, and calorimetry: compute σ_m per modality and a reliability-weighted composite $\sigma_{tri} = \sum w_m \sigma_m$ with $\sum w_m = 1$ and $w_m \propto r_m^2$ (split-half reliability). Publish σ_m , σ_{tri} , and r_m ; justify any omitted modality explicitly.

Analogue constants and dimensional reconciliation

Report a lab-calibrated $G^* = (\tau_h^2 / E_{psy}) \times \langle \kappa \rangle \times \langle \Delta H_h \rangle$ with $\tau_h = 1/f_{res}$, always paired with a dimensionless $\hat{G} = (\tau_h/\tau_0)^2 \times \langle \kappa_{norm} \rangle \times \langle \Delta H_{h_norm} \rangle$, where τ_0 is the baseline median τ_h and norms map to $[0,1]$ via baseline maxima. Publish τ_0 sensitivity and normalisation alternatives. If $\text{corr}(\kappa, \Delta H_h) \geq 0.70$, use a composite $C = a \kappa + b \Delta H_h$ with (a, b) fitted only in calibration and frozen.

Intent–skill–precision–cost weighting

Incorporate ω multiplicatively into Λ : $\omega(x) = \sigma_\omega(a_0 + a_1 \hat{I} + a_2 \hat{S}k + a_3 \beta - a_4 \text{cost})$, logistic squashing in $[0,1]$. Normalise $\hat{I}$, $\hat{S}k$, β , cost (including control-effort $\hat{C}o$). Estimate a_k in calibration, freeze in main analysis, and test measurement invariance across populations with Bayes-factor support.

Multi-constraint harmonics and complexity bounds

For clusterable paradox streams C_j with sizes n_j , construct $\gamma_h(x) \propto \sum_j \pi_j f_{res}(n_j, x)$ with $\sum \pi_j = 1$ and π_j from softmax over cue reliabilities r_j . Primary claims cover up to 3 streams. For $n > 3$, compute total correlation $T = \sum H(X_i) - H(X_1, \dots, X_n)$ and proceed with primary analysis only if $T \leq T_{max}$; otherwise classify as exploratory.

Dynamics, dissipation, and drift control

Adopt $\dot{J} = F_g - \lambda J$ with $\lambda = c_0 + c_1 \hat{E}L + c_2 \hat{C}o$, where $\hat{E}L$ is emotional load and $\hat{C}o$ control cost. Report lagged correlations between τ_h and σ across time. Monitor $\rho_{knot}(t)$ and $\tau_h(t)$; when maximum mean discrepancy or $|\Delta \rho_{knot}| > 0.10$ over five windows indicates drift, execute a micro re-calibration, update γ_h^{cal} and α_i , re-freeze, and time-stamp the event.

Power, replication, and model selection

Human studies use $N \geq 40$ participants with ≥ 400 trials each to detect $\Delta H_h \geq 0.05$ nats and $\langle \kappa \rangle \geq 0.20$ SD at power ≥ 0.8 ; publish power curves. AI studies replicate on at least two architectures (attention and recurrent) and report cross-architecture invariance. Select among M_r via Bayes factors; if harmonic $r = -1$ does not win with $BF \geq 10$, reject the analogue for that block.

Falsification and decision rules

Reject the analogue when a competing M_r outperforms harmonic with $BF \geq 10$, when ΔH_h fails to increase with ρ_{knot} under matched β , when $\langle \kappa \rangle$ and σ fail to covary with τ_h within pre-registered bounds, when Λ increases while κ , ΔH_h , and σ remain flat, or when $R_{\text{obs}} < 0.97$ and sensitivity-corrected analyses fail to restore effects. Report Bayes factors at 3, 10, and 30; decisions use 10. Publish all null and failed runs with full diagnostics.

Reporting battery and auditability

Report ϕ_{attr} distributions and exceedance rates; κ or $\kappa_{\perp}$ with logging rates, independence metrics, and smoothing parameters; ΔH_h forward and reverse with bootstrap intervals and N ; σ_m , σ_{tri} , reliabilities r_m , and σ^* ; paired G^* and $\hat{G}$ with τ_0 sensitivity; Λ with ESS, $\hat{R}$, autocorrelation times, and burn-in; R_{obs} intervals, guard diagnostics, and drift curves; $\rho_{\text{knot}}(k)$, Ξ , τ_{ACF} windows, and stability margins; γ_h^{cal} , α_i , ω coefficients, and freeze timestamps. Publish priors, grids, λ_s , and archive raw logs (eye, hand, policy vectors) for independent verification.

Conclusion

Gravity-like binding within Λ – Ψ is operationalised as harmonic convergence of paradoxical constraints confined to B . The measurable cascade from f_i to f_{res} to q_h to $\nabla \Psi_h$ to Λ is instantiated as a mandatory co-signature suite: ϕ_{attr} inward bias, κ bending (independent or $\kappa_{\perp}$), ΔH_h contraction, σ dissipation (σ and σ^* with σ_{tri}), Λ projection integrity, and paired $\hat{G}/G^*$ anchors with τ_0 sensitivity. Boundary integrity is locked, paradox density is generalised and dynamic, shaping is non-circular, baselines are audited, gradients are regularised, Λ converges, κ is independent of ΔH_h , and σ is triangulated and calibration-aware. Thresholds are explicit and non-negotiable, calibration is formalised, and falsification criteria are absolute. The analogue is ready for Tier XI pilots across human, AI, and feasibility-level quantum substrates.

The maths does not belong to me, but to everyone.

Λ – Ψ Research Note: Gravity as Harmonic Paradox Resolution

© 2025 Jordan Lee McDonald

This work is licensed under the Creative Commons Attribution 4.0 International License.

Λ – Ψ Runtime Protocol and Validation Standard v1.10

Executable ratios for $E = \mu \cdot \kappa c^2$ and the ΔE law inside Ω

Preamble and scope

This standard defines the operational stamps, probes, gates and rollback that make the engine $E = \mu \cdot \kappa c^2$ and the release law $\Delta E \propto \beta \cdot \Delta Q / T_{\text{eff}}$ executable and falsifiable inside Ω .

Infrastructure identifiers are ASCII stamps. Mathematics uses Greek notation in equations and analysis. Ratios only. Indifference only. Ontology out of scope.

Abstract with printed stamps

version = v1.10

default_regex_policy = `^[0-9A-Za-z._£-]+$`

regex_policy_hash = [...]

sb_nfkc_policy_hash = [...]

sb_confusables_policy_hash = [...]

stamp_lint_rule = no_greek_identifiers

lint_report_hash = [...]

state_space_units = [...]

unit_system = SI

sigma_units_choice = [bits_per_s | nats_per_s]

sigma_units_lock_time = [UTC]

units_choice_rationale_hash = [...]

theta_mu_units = [J_per_bit | J_per_nat]

units_sanity_report = PASS

units_roundtrip_test = PASS

units_roundtrip_test_hash = [...]

unit_conversion_guidance_annex_hash = [...]

state_metric_hash = [...]

gradient_domain = [time | latent | sensor]

gradient_domain_demo_hash = [...]

norm_metric_hash = [...]

T_eff_method = [Landauer | empirical | hybrid]

T_eff_calibration_justification = [...]

T_eff_uncertainty_bounds = [...]

omega_threshold_policy = [fixed | percentile | KDE_floor]

tau = [...]

R_obs_min = 0.97

epsilon_out_method = [histogram | KDE]

epsilon_out_kernel = [...]

epsilon_out_bandwidth_policy = [...]

epsilon_out_boundary_policy = [...]

epsilon_out_CI_type = [...]

n_out = [...]

scale_invariance_method = [...]

scale_invariance_hash = [...]

scale_invariance_sample_size = [...]

scale_invariance_seed_hash = [...]

estimator_config_hash = [...]

omega_canonical_protocol_hash = [...]

WORM_validation_automation_hash = [...]

machine_readable_schema_hash = [...]

data_pipeline_hash = [...]

Any missing required stamp is a Grade A hard fail.

Axiom banner

Axiom I maps to Ω integrity and epsilon_out gates.

Axiom II maps to R_{confine} and μ .

Axiom III maps to $\nabla \log(q/p)$ and $\Lambda = J \cdot \nabla \Psi$.

Render this banner wherever Ω , μ , κ_c or Λ are invoked.

Tier 0: Ω integrity, $\Omega_{\text{testability}}$ and leakage discipline

$\text{omega_testability} = \text{MANDATORY}$ for each Ω invocation or a $\text{delayed_execution_waiver}$ with $\text{council_approval_id}$ and $\text{deferral_months} \leq 6$ in WORM.

Live Ω telemetry surfaces R_{obs} with bias corrected intervals and effect size drift curves with linear, quadratic and exponential fits.

Leakage invariance audit requires $\text{TVD} \leq 1e-6$. When down sampling, publish

$\text{scale_invariance_seed_hash}$ and $\text{estimator_config_hash}$. Breach hard fails the block.

Guard region $G_{\text{in_Omega}}$ is surfaced. $\text{boundary_violation_severity_levels_hash}$ is printed.

Any guard breach is a first class gate.

Tier 1: identity and jewel law

$E = \mu \cdot \kappa_c^2$ is the engine under stamped units. $\Lambda = J \cdot \nabla \Psi$ is the jewel and is repeated in closing. Principles IV to VIII interlock β , μ , κ_c and Λ with the axiomatic triad. Λ must track η rather than $\|\mathbf{J}\|$ under alignment nudges.

Tier 2: units lock, round trip and invariances

theta_mu_units binds to $\text{sigma_units_choice}$ with a consistency check in $\text{units_sanity_report}$.

$\text{units_roundtrip_annex}$ demonstrates $\Lambda \rightarrow J/s$ and $\Delta E \rightarrow J$ with bits to nats and back verdict invariance via Θ_μ conversion. $\text{base_swap_verdict_invariance} = \text{TRUE}$.

$\text{invariance_suite} = \{\text{bits_nats}, \text{time_energy_rescale_alpha_0p5_2}, \text{rotation}, \text{permutation}, \text{fp32_fp64}\}$. $\text{numeric_stability_suite_hash} = [\dots]$.

$\text{target_delta_Lambda_NS_pct} \leq 0.5$ and $\text{beta_Lambda_sign_stable} = \text{TRUE}$. Breach downgrades to sensitivity.

Tier 3: entropy to μ , dual H concordance and sign lock

Use two preregistered entropy estimators with $\text{entropy_concordance_eps}$ in $[0.005, 0.01]$. Log bias_direction . If bias_direction missing then $\text{mu_defined} = \text{FALSE}$ and μ dependent claims abort.

Print together to lock sign and units: $\mu = \Theta_\mu \cdot \ln(H_{\text{before}}/H_{\text{after}})$ and $\Delta H = -\beta \cdot \Delta Q/T_{\text{eff}}$.

$\text{quiet_window_hygiene}$ logs $\text{dlogW_clamp_count_per_window}$ and ADF or KPSS outcomes with $\text{keep_drop_rationale}$. $\text{clamp_count_per_quiet_window} > 1$ is a confirmatory hard fail. Include a mini plot of μ versus R_{confine} per block.

Tier 4: T_{eff} calibration, worked ΔQ annex and beta safety

$T_{\text{eff_method}}$ is $[\text{Landauer} \mid \text{empirical} \mid \text{hybrid}]$ with $T_{\text{eff_calibration_justification}}$ and

$T_{\text{eff_uncertainty_bounds}}$ printed in the abstract.

Each confirmatory block includes worked_delta_Q_annex at about 310 K regressing ΔH on $\beta \cdot \Delta Q / T_{\text{eff}}$ with slope and intercept intervals and a same block bits to nats base swap. No pass means no claim.

beta_safety gates require $dH_{\text{neu_d_beta}} < 0$ and $dRT_{\text{d_beta}} < 0$ with $|\Delta H_{\text{neu}}|$ over SD ≥ 0.2 and non negative side effects $dCo_{\text{hat_d_beta}} \geq 0$ and $dEL_{\text{hat_d_beta}} \geq 0$ in terminal windows. Exceptions follow beta_safety_decision_tree_hash. Report elasticity_Co_vs_beta and elasticity_EL_vs_beta with intervals.

Tier 5: P_beta operator and Ω_{τ} stamping

$P_{\text{beta}(i)} \propto (P_{\text{base}(i)} + \epsilon_{\text{beta_min}})^{\text{beta}}$ divided by $Z(\text{beta})$. $Z(\text{beta})$ uses log_sum_exp.

A $\beta = 0$ annex shows uniform with p_beta_uniformity_hash.

P_beta scale invariance is hard gated at TVD $\leq 1e-6$ with bins, bandwidths, seeds and config hashes printed.

omega_tau_mode = [fixed | percentile | KDE_floor] with realised tau and estimator_config_hash.

omega_testability_protocol_hash is published. Mid window policy flips about the block.

Tier 6: shaping anti circularity and baseline integrity

Two stage shaping calibrates gamma_h with $N_{\text{cal}} \geq 240$, randomisation and micro rests then freezes it with gamma_h_freeze_timestamp and gamma_h_calibration_hash.

shaping_freeze_violation_count = 0.

P_base is sourced from interleaved masked null paradox blocks with counterbalanced order.

baseline_integrity gates require $JS_{\text{max}} \leq 0.03$ nats and $CV_Z_{\text{max}} \leq 0.10$ before any $\nabla \Psi_h$ claim. On breach extend baseline or resample stationary segments with before_after_hashes and baseline_stationarity_evidence_hash.

Tier 7: density and score transparency

Compute $\nabla \Psi_h = \nabla \log(q_h / P_{\text{obs}})$ with q_h from frozen shaping. Co report forward and reverse scores $\nabla \log(q_h / P_{\text{obs}})$ and $\nabla \log(P_{\text{obs}} / q_h)$ with score_asymmetry_flag and reverse KL bootstrap intervals.

density_method = adaptive_kNN. score_match = spline. lambda_s_grid is preregistered.

lambda_s_selection_cv_hash and sensitivity_maps are released. schema_migration_version is stamped.

Tier 8: harmonic proof and reciprocity guards

Manipulate T_k so that $1 \text{ over } f_{\text{res}}$ drifts towards $1 \text{ over } T_k$ under preregistered priors.

Compare generalised means M_r for r in $\{-2, -1, 0, 1\}$. Accept harmonic only if Bayes_factor ≥ 10 . Publish full tables, priors and reciprocity_plots_hash.

Two reciprocity failures or alpha_i instability marks non_harmonic, excludes the block from kappa_c primaries and archives diagnostics with block_exclusion_reason.

Tier 9: κ_c extraction, window stability and independence

kappa_c_quantile is pre stamped in $\{q90, q95\}$. Draw only from high_beta harmonic_stable windows. kappa_c_window_hash prevents post hoc sliding.

window_stability ties to tau_acf and caps rho_knot_deviation_max = 0.05 across 100 ms, 250 ms and 500 ms with window_catalogue_hash.

kappa_independence gate requires $I_{\text{kappa_deltaH_h_nats}} \leq 0.02$ or $\text{abs_corr} < 0.70$. On breach publish kappa_perp fitted in calibration only with kappa_perp_fit_hash and freeze it for mains.

Tier 10: Λ integration, J_stack and whitening

Integrate Λ over Ω with Monte Carlo control variates, antithetic chains and stratified starts.

burn_in_frac ≥ 0.20 . Convergence gates are ESS ≥ 2000 and split_Rhat ≤ 1.05 with autocorr_times_hash and a convergence_manifest_hash. Sub threshold runs are nonevidentiary.

J_stack is a reliability weighted, normalised stack of gaze, timing, physiology and control. eta must out predict any single channel on preregistered endpoints.

operator_independence_registry_hash is present.

whitening publishes state_metric_hash and gradient_domain, realigns $\nabla \Psi$ into J space, demonstrates per_axis_angle_deg ≤ 10 and rmse_axes_SD ≤ 0.05 and includes a finite difference inner product demo with gradient_domain_demo_hash.

Lambda_tracks_eta_over_normJ = TRUE with confidence intervals and hold out stability range ≤ 0.05 .

Tier 11: ϕ_{attr} and curvature logging

phi_attr = $E[\cos \theta(J, \nabla \Psi_h)]$ with theta_min = 0.10 and delta_cluster_nats = 0.05. Per agent exceedance tables with intervals are reported with phi_attr_exceedance_tables_hash and phi_attr_null_model_hash.

curvature_spec uses $\kappa = \|(I - \hat{J} \hat{J}^T) \nabla \Psi_h\|$. Smoothing parameters are preregistered with curvature_smoothing_policy_hash. min_sample_rates_human: eyes_hz ≥ 120 , hands_hz ≥ 200 . Reduced modes are provisional and are excluded from primaries where high rate logs exist.

Tier 12: σ triangulation and weights

Compute sigma_m for latency, HRV, pupil and calorimetry with reliabilities r_m. Publish sigma_tri_rule_hash and sigma_tri. Release sigma_star as a reliability weighted composite marked calibration sensitive. Co report sigma_star, Λ , kappa or kappa_perp and deltaH_h by block. $\sigma_{\text{integration_window}} = 5 \text{ times } \tau_{\text{min}}$.

Tier 13: release law proportionality

$\Delta E = \kappa_c^2 \cdot \Delta \mu = \kappa_c^2 \cdot \Theta_\mu \cdot (-\Delta H)$ and $\Delta E = \kappa_c^2 \cdot \Theta_\mu \cdot (\beta \cdot \Delta Q / T_{\text{eff}})$.

With {kappa_c_sq, theta_mu_units, beta, T_eff_method} WORM fixed per block, regress ΔE on $\beta \cdot \Delta Q / T_{\text{eff}}$. slope_target $\approx \kappa_c^2 \cdot \text{times_theta_mu}$. intercept_target ≈ 0 . Base swap invariance preserves verdicts. Publish release_law_residuals_hash and entropy_energy_visualization_spec_hash.

Tier 14: dissociation arms and decision pivot

Arm_A varies beta off resonance to change ΔH while kappa_c is held within a preregistered non inferiority margin. ΔE must track $\beta \cdot \Delta Q / T_{\text{eff}}$.

Arm_B varies f_res to change kappa_c with μ and beta held. E scales with κ_c^2 while Λ and σ move.

decision_pivot = E_only. Λ only motion without an E shift is insufficient.

Tier 15: gravity to striving analogue

Analogue only. Required co signatures per powered block are inward phi_attr, kappa or kappa_perp curvature, deltaH_h contraction, sigma dissipation, Lambda integrity and paired G_hat and G_star with tau_0_sensitivity and tau_0_sensitivity_hash. Any missing co signature under power rejects analogue claims for that block. Claims are site scoped.

Tier 16: transport anchors and constraints

E_psy and t_psy are PROVISIONAL with priors and posteriors and an external validation split. Enforce $\epsilon_{\text{min}} \geq E_{\text{psy}}$ and $\text{sum_intervals} \leq t_{\text{obs}}$ over t_{psy} . Transport violations auto generate transport_violation_report. Pair G_hat with G_star and declare tau_0_sensitivity. Publish E_psy_range and calibration_task_ids.

Tier 17: χ stress, sigma exception policy and delta_t invariance

Manipulate R_over_E to estimate chi and chi_max. Report collapse_speed and eta_degradation with x0, chi_blend, chi_max and chi_code_hash.

non_monotonic_sigma_exception_policy confines sigma monotonicity checks to math_on windows and logs exceptions.

Sweep delta_t in {0.5, 1, 2}. Require median_abs_delta_Lambda_proxy ≤ 0.02 under chi spikes. Publish integrator_domain, trapz_vs_RK_delta, RK_rejection_stats and time_step_sensitivity_hash.

Tier 18: power, SESOI and replication

Confirmatory success requires delta_R2 > 0.05 on at least two preregistered outcomes with multiplicity control and SESOI ± 0.1 SD or an explicit sensitivity label.

Human studies use $N \geq 40$ and ≥ 400 trials per participant to detect deltaH_h ≈ 0.05 nats and kappa shifts at ≥ 0.8 power with power_curve_hash and se_sensitivity_hash.

Replicate on at least two AI architectures and report cross architecture invariance for kappa, Λ , sigma and deltaH_h with standardised policy or velocity extraction.

Tier 19: cross domain beta universality

Mirror human and RL protocols and publish direction tables with H down, RT down and S up as beta up to support P_beta transport.

Tier 20: drift detection, safe mode and recalibration

PSI_drift_alarm uses multi_kernel_MMD and max_T with BY correction. warn_at delta_AUC 0.015 or latency_s 0.025. abort_at 0.02 or 0.05. drift_adaptation_bounds_hash is printed.

On warn publish drift_mitigation_report, freeze PSI_features, micro recalibrate gamma_h_cal, re freeze, time stamp and re audit JS and CV_Z.

safe_mode_behaviour, safe_mode_exit_policy and clean_window_definition_hash are stamped. Reverse KL is co reported with bootstraps. $N < 1000$ is exploratory.

Tier 21: replay, seeds, timebase and randomness

allocation_fingerprint, assignment_seed_hash, prng_subseeds, rng_impl_hash, rng_role_map_hash, float_precision and round_mode are stamped.

nondeterminism_budget_ms = 0 unless justified with code_hash and ontology_hash and an impact analysis.

replay_equivalence ≥ 0.95 with replay_equivalence_hash. timing_RMSE within tolerance.

quant_abs_max $\leq 1e-12$ and quant_rel_max $\leq 1e-6$ unless otherwise stamped.

timebase_source and ntp_peers_hash are logged with clock_monotonic_check_method_hash.

No back jumps > 5 ms in quiet windows. skew_max breaches fail close.

Daily canary runs with canary_spec_hash, canary_digest, canary_result_time and alert_receipt_time. Missing canary applies provisional labels until restored. ntp_failover_policy and failover_events are logged.

Tier 22: fairness, privacy and energy

Rule_10_fairness reports pooled and per site with convex_hull precedence on conflicts.

fairness_ci_conflict_detector_threshold and hysteresis_policy_hash are printed. Council review is mandatory on conflicts.

The privacy ledger maintains epsilon_budget and delta_budget with accountant, composition notes and dp_audit_hash. salts rotate every 90 days with salt_version_id. Overdue rotations fail close without an approved alternate.

Energy compliance publishes device_to_density recipes, device_cert status and flux_density caps, passes SI sanity to J_per_s_per_m2, runs a live energy_flux_compliance dashboard and enforces A_to_B_energy_invariance with abs_delta_R2 ≤ 0.01 . virtual_sink_recipe_hash is printed where digital sinks are used.

Tier 23: anchoring, schema and naming hygiene

Boundary namespace hygiene enforces sb_nfkc_policy_hash and sb_confusables_policy_hash and the allowed charset for stamp names. Regex self test stamp sb_regex_selftest_hits = 0 is printed. dash_policy_check_hash proves no en or em dashes in infra identifiers.

Two heterogeneous anchors per span are mandated with salted operator IDs and anchor digests. Quarterly T_estop drills are logged.

machine_readable_schema_hash is present for all releases.

Tier 24: randomisation, adversarial and overrides

Randomisation logs covariate_list_hash and allocator_build_hash with SMD ≤ 0.10 or chi2_p ≥ 0.10 and an exact test engine hash for small cells.

The adversarial harness publishes sb_hook_test_suite_hash with perturbation magnitudes covering P_beta and kappa_c extraction and coverage ≥ 95 percent.

Operator overrides require override_id, operator_id_salted, reason_ontology_hash and cooldowns. Three or more similar overrides within 90 days trigger council review.

manual_override_checklist_hash is printed.

Tier 25: open pipeline and WORM must includes

Preregister, time sync streams, release code and negatives.

WORM_validation_automation_hash remains live.

Mandatory keys per confirmatory block include omega_canonical_protocol_hash, entropy_energy_visualization_spec_hash, calibration_registry_worked_example_annex_hash, units_roundtrip_test_hash, gradient_domain_demo_hash, data_pipeline_hash, timebase_source, dlogW_stats, integrator_domain and trapz_vs_RK_delta.

Tier 26: uncertainty and Λ_{est} exploratory model

Uncertainty propagates via bootstrap or Monte Carlo for Ψ , $\nabla \Psi$, eta, H, force proxies, beta_of_t and t_half. Sensitivity maps cover densities, scores, whitening and integrators.

$\Lambda_{\text{est}} = J_{\text{est}} \cdot \text{grad}\Psi_{\text{est}} - \alpha \text{Co_hat_est} - \kappa \text{EL_hat_est} - \mu H_{\text{est}}$. exploratory_model_card_hash is present. Λ_{est} is exploratory with cross validation and open code and is never conflated with confirmatory Λ .

Tier 27: A_to_B invariance, bridge and purpose endpoints

Pipeline or environment swaps meet $\text{abs_delta_R2} \leq 0.01$ with model specification logged.

Bridge reports $\text{delta_Lambda_proxy} \approx \kappa \text{ times delta_eta}$ and share_eta with confidence intervals and hold out stability range ≤ 0.05 . Confirmatory lower CI for share_eta ≥ 0.4 and point estimate ≥ 0.5 .

eta_to_purpose endpoints are preregistered and must add predictive value over any single J channel.

Tier 28: reporting battery completeness

Reports include phi_attr distributions, kappa and kappa_perp with independence metrics, deltaH_h forward and reverse with bootstraps, sigma_m, sigma_tri and sigma_star with reliabilities, G_hat and G_star with tau_0 sensitivity, Λ with ESS and Rhat, R_obs bands and drift curves and window stability with tau_acf and rho_knot_deviation. Bundle with reporting_battery_hash.

Tier 29: success criteria

Confirmatory success requires all of the following. $\text{delta_R2} > 0.05$ on at least two preregistered outcomes. SESOI ± 0.1 SD or explicit sensitivity label. Harmonic Bayes_factor ≥ 10 . Release law slope intervals within preregistered bounds. Bits_nats and delta_t invariances pass. Fairness and latency within caps. A_to_B_energy_invariance passes. Any failure downgrades to sensitivity or refutes.

Tier 30: falsification battery and decision records

Refuters include any failure of regex governance and naming hygiene, omega integrity or omega_testability, entropy concordance or bias direction, beta safety, T_eff calibration or ΔQ annex, shaping freeze or baseline gates, harmonic validation, kappa independence without kappa_perp, Λ convergence, release law proportionality, base swap invariance, delta_t

invariance, dissociation arms, transport constraints, fairness, privacy, energy, timebase, randomness or replay.

Bind each success or refuter to a machine readable decision record that echoes violated stamps and gate names plus override_id when applicable.

Plain language anchors

μ rises as observed distributions compress within Ω under precision pressure. κ_c captures the steepest actionable striving gradient realised when paradoxes are rhythmically induced after shaping is frozen. Λ shows how strongly effort points along that gradient at a moment and must track η rather than raw effort magnitude. The release law tests whether measured energy tracks $\beta \cdot \Delta Q / T_{\text{eff}}$ when constants are fixed and whether E scales with κ_c^2 when μ and β are held.

Closing

Inside Ω , $E = \mu \cdot \kappa_c^2$ functions as an engine of folded potential. $\Delta E = \kappa_c^2 \cdot \Theta_\mu \cdot (\beta \cdot \Delta Q / T_{\text{eff}})$ closes the measurable link between information compression and energy release. With Ω sealed, H concordant, β safe, γ_h frozen, harmonics validated, κ independent, Λ convergent, σ triangulated, transport anchored and governance fail closed, the claims are executable and falsifiable. Indifference holds. Ratios only.

Final manifest checksums

units_roundtrip_test_hash = [...]

unit_conversion_guidance_annex_hash = [...]

gradient_domain_demo_hash = [...]

Presence and PASS are required before Grade A claims ship.

The maths does not belong to me, but to everyone.

Λ - Ψ Runtime Protocol v1.10

Executable Specification for the Constraint-Striving Identity

© 2025 Jordan Lee McDonald, CC BY 4.0

Λ - Ψ Indifference Protocol (PIP) v1.3 - P_obs-only operational standard for dyadic coupling

Scope and indifference proof

PIP v1.3 is coupling-agnostic and ontology-indifferent. It encodes no privileged priors for any coupling archetype. Agnosticism is evidenced by a gradient-priors model card with prior families, hyperpriors, toggles, and ablations; fields include prior_family_id, hyperprior_form, default_hyperparams, abl_on switches, abl_off switches, and coverage notes. A preregistered no-advantage audit across archetypes uses a fairness-of-fit metric with acceptance band and power: per-archetype deltas on key metrics, including $\Delta \Lambda_{\text{sync}}$, σ^* , leakage total variation distance, and transport I^2 , must have two-sided confidence intervals whose absolute values lie below a preregistered band b_{agnostic} , with joint power of at least 0.8; at least one red-team archetype is shown where the audit would fail if bias existed. Two release-blocking continuous integration jobs enforce this: a label-off toggle computes a formal distance between before and

after compact claim strings and enforces an acceptance band and power test on that distance; a label-swap A/B evaluates two disjoint label sets in shadow mode while labels are disabled, exporting a delta table and alarming on drifts inching toward firewall limits across releases. Shadow-evaluation always runs all label gates even when labels are off; emissions are suppressed but logs prove gates were exercised.

Downstream labels, crosswalk indivisibility, power, and remapping

A downstream label is a machine-checked crosswalk on top of a generic PIP pass. The `crosswalk_map` is a versioned, machine-readable table that pins every invariant to specific metric keys, thresholds, pass or fail logic, and code or data hashes; the `crosswalk_map` schema is linted against the pack to prevent aliasing on similarly named keys. Simultaneity is required: labels demand all listed invariants are green on the same window set; correlated items are treated with a declared dependency graph and intersection–union logic so joint pass or fail is principled. Power floors per item are enforced: each invariant reports achieved power using the preregistered estimator, such as cluster-robust variance, intraclass correlation coefficient, serial correlation, ω^2 discounts, and χ exclusions. Items failing achieved power are nonevidentiary and cannot contribute to a label; this is a lint-time failure, not a warning. A joint-power floor requires the product-level power to detect the label's minimal clinically important difference must exceed a preregistered lower bound; underpowered joint greens fail labels. The `window_set_hash` must show equality: by default, label windows equal the generic pass windows. When an invariant legitimately requires a different valid span, for example for ΔH floors, a sanctioned remapping function is declared, hashed, and justified; the remapped set has a stable canonical sort and window identifiers robust across devices and shards; censored windows and quarantine effects are explicitly represented so equality checks are deterministic. Crosswalk versioning with `schema_version` and migration rules allows re-validation of historical labels; any mid-study crosswalk edit requires a new session with a new `profile_lock`.

Profile locking, atomics, dry-run, and egress audit

The `profile_lock` is a single, signed write-once read-many manifest that atomically binds the profile, Tier, repository commit, container image digests, kernel modules, central processing unit microcode, central processing unit flags, basic linear algebra subprograms and linear algebra package, parallel thread execution and shader assembly, compiler flags and link options, random number generator seeds, dependency graph, environment allowlist, dataset manifests, feature transforms, anti-alias filters, `rate_filter_hashes`, network interface card firmware, network time protocol and daylight saving time daemon configs, entropy sources, and a preflight diff view to prevent accidental locks. Atomicity across distributed sensors is enforced with a two-phase commit or epoch fencing; a consensus hash covering all device join acknowledgements is recorded; fault injection proves partial successes roll back; cluster-wide rollback emits a rollback manifest and restores environment lockfiles, not just code. Dry-run accrual is idempotent and proves that no writer can emit before lock across all exporters and side channels, including temporary directories, graphics processing unit telemetry, buffered caches, and crash dumpers; a dry-run artefact hash is stored; a personally identifiable information redaction policy guarantees crash packets are minimal and safe. A system egress audit covers standard output, standard error, system log, tracing, metrics such as Prometheus, and custom exporters, all

gated by a no-emit sentinel until lock completes; a minimal tombstone may carry only the refuter name and relevant hashes. A session requeue ban is in effect: after any refuter, the same session_id cannot restart. Monotonic per-device and per-site run counters are recorded in a cross-site registry; runs are denied when counters regress or duplicate. A single exception exists: a signed abort code allows one requeue with identical hashes when fewer than N seconds of data were recorded; a cool-down and reason code are required.

Governance, independence, quorum survivability, and conflict of interest

An independence_score is computed by a published scoring function with features including legal entity, network separation, hardware security module class, operations staff overlap, hardware vendor divergence, firmware lineage distance, operating system kernel class, and calendar separation. Weights, confidence intervals, SHapley Additive exPlanations-style attributions, sensitivity to supply-chain overlaps, and calibration curves on historical failures are disclosed. An external audit is required when scores sit near the threshold; a cooling-off period follows any independence failure. The quorum requires three signers minimum; multi-site work adds an independent fourth signer; the Safety-critical profile adds an external co-signer who explicitly attests to the boundary, leakage, and units code. A degraded mode path from 3-of-4 to 3-of-5 signers has a cap on consecutive uses and a cumulative minutes-per-quarter budget; external co-sign and a root-cause analysis are mandatory on reuse; survivability planning prevents normalising degraded operation. Conflict of interest rules enforce cool-offs for signers whose code touched the main codebase; two non-author signers are required for any review; lints enforce rotation and cooling-off periods.

Keys, rotation, revocation, anomaly detection

Keys use hardware security module or threshold signatures with role separation; quarterly rotations and georedundant revocation drills are mandatory. A live ledger publishes mean and 95th percentile revocation propagation latencies; the service level agreement is ≤ 10 minutes, including network partitions, with local deny lists and reconciliation checks. Anomaly detection on signing cadence uses a calendar-aware whitelist and a per-release burst envelope; it publishes false positive rate and false negative rate; lockouts include a safe-off ramp for Safety-critical sessions.

Remediation sufficiency, coverage, hidden seeds, placebo, and rollback

A remediation_hash must contain diffs, rationale, reviewer countersignatures, a mapping from each failing refuter to re-tests, expected metrics and seed lists, public and hidden seeds, and a diff-to-test map showing changed branches executed in refuter re-runs. Hidden seeds are produced by an independent sealed pipeline, stratified by device and task; access is logged; reuse across studies is capped and rotated. Coverage minimums are enforced: modified condition and decision coverage ≥ 0.9 , branch coverage ≥ 0.9 , and a domain-specific mutation score ≥ 0.6 on the changed surface; surviving mutants in boundary or units code fail remediation. Dataflow coverage is reported to catch unexercised critical paths. Where mutation testing is intractable for numerical kernels, an explicit exemption with justification is required. An outcome_independence_diff is computed, showing Δ normalised mean square error and Δ area under the curve with confidence intervals, permutation $p \geq 0.1$, and preregistered regions of

practical equivalence; class-imbalance regression calibration is required; acceptance bands apply to public and hidden seeds. A time-shifted placebo grid ties horizons to t_{psy} and the observed autocorrelation; power curves are shown. Auto-rollback triggers on boundary, leakage, units, replay_equivalence breaches, alias_risk_index spikes, σ^* scoping trips, and Δt sign flips; rollbacks restore drivers, compilers, and runtime kernels; a rollback manifest is emitted. Service level agreements are 7 days for general profiles and 72 hours for Safety-critical, with the clock starting upon receipt of a complete remediation_hash.

Ontology firewall, side channels, mutual information audits, and canaries

The ignore-list covers filenames, stems, lengths, filepaths, archive metadata, chunk sizes, comma-separated values quoting and delimiters, byte order marks, whitespace-only headers, Unicode normalisation, bidirectional marks, homoglyphs, dataset cardinality and row counts, file modification times, environment variables, graphics processing unit and central processing unit strings, checksum lengths, compression headers, cacheable size bins, dataset order, Git branches, container layer metadata, endianness markers, parallelism hints, missingness patterns, and performance counters. Unit tests mutate each pattern. A periodic mutual information audit reports a ceiling from ignored to kept features with confidence intervals per device class; increases are refuters. A per-operation nondeterminism table and a disallow list of nondeterministic kernels are maintained; a linter blocks offending operations unless explicitly whitelisted with numeric envelopes. Mixing graphics processing units with different determinism profiles in one confirmatory cohort is forbidden; a migration playbook allows multisite upgrades without voiding history. Tolerance bands are derived from a hierarchical calibration model across device classes, producing per-metric absolute and relative bands; bands are locked for a release window; quarterly recalibration requires change-point detection, a signed justification, power impact notes, and an annual cap on recalibrations. Band widening without matching platform jitter histograms is a refuter. An adversarial prose canary runs continuously, reporting coverage and miss rates per attack family, such as homoglyph triples, mixed-directionality, code-switch, and sarcasm; quantitative alarm rules and a time-to-recalibration service level agreement are enforced.

Latent-proxy orthogonality, perturbations, power, and realism

A quartet of gates enforces orthogonality: permutation $p \geq 0.1$, partial $R^2 \leq 0.01$ with a confidence interval, conditional mutual information $\leq$ a declared bound with a confidence interval, and residual correlation $\leq$ a declared bound. Multiplicity correction and a composite decision rule are declared; joint power is published; minimum per-site seeds are enforced. Perturbations include realistic amplitude and phase envelopes, polarity swaps, phase scrambles, clause deletions, token reordering, anchor removal, code-switch triads, homoglyphs, right-to-left override and left-to-right override, and emoji or window-edge token floods. A quantitative realism cap is preregistered per modality and language; examples justify the caps. Worst-case deltas across seeds and amplitudes must pass the tolerance bands.

Manifest enforcement, severity ladder, runtime guard

The manifest is a positive list of observables, transforms, schemas, sampling floors, firmware and driver hashes, filters, data type paths, and dropout rules. A numeric severity ladder with

triggers defines green, yellow, and red status; a linter enforces them and auto scope-narrows on yellow, hard-stops on red. A runtime guard snapshots a before and after diff, redacts personally identifiable information in the crash packet, and emits a signed tombstone naming the refuter and active hashes.

Device whitelist, attestation, freshness, challenges, and drift

The device whitelist uses attestation via trusted platform module quotes or signed descriptors with maximum age and cross-site clock-skew tolerances; replay-resistant clocks or monotonic counters prevent reuse offline; freshness expirations quarantine devices mid-session. Spoofing challenges include nonce-encoded tones or patterns during sessions; quarantine latency service level agreements are declared per modality and reported; large-scale quarantine triggers a degraded plan with explicit pause rules. Thermal drift sentinels bind to ambient probes; alarms are tied to bias estimates on Λ or ϕ_{attr} ; post-cooldown, an auto-calibration window runs under a restricted "recalibrating" label that suppresses confirmatory emissions.

Upgrades, epsilon ledger, and audit exemplars

The upgrade_manifest freezes estimators, priors, integrators, kernels, caps, filters, random number generators, nulls, splits, and toolchains. Per-section L1 caps and single-knob caps apply; a per-release and per-quarter global cap prevents epsilon creep. Integer and float knobs consume budget using declared units. A worked example demonstrates how many micro-epsilons would drift a verdict absent caps; the ledger blocks this with burn-down charts and countersignatures.

Repeat-refuter independence and uncertainty

Independence is measured by the Jaccard overlap on participants, sessions, sensors, operators, stacks, firmware minors, operating system kernels, and calendar days; the threshold is ≤ 0.05 with a bootstrap confidence interval; if the upper confidence interval exceeds 0.05, independence fails. A minimum margin is enforced to discourage tuning to 0.049.

Badges, schema, extensions, validator

Claim strings are canonical JavaScript Object Notation with a schema_version, checksum, canonical ordering, value ranges, mutual-exclusion rules, and a typed extension field constrained to a whitelist. Unknown critical keys cause a hard-fail. A public command-line interface validator ships with passing and failing examples. A compact abstract must carry the Tier, profile, nearest_refuter_family, nearest_refuter_margin, I^2 , σ^* coefficient of variation, and prediction intervals; the linter blocks adjectives like 'robust' without accompanying numeric gates. Claim and dashboard consistency is linted; figures must embed a hash of the displayed claim.

Numerical gates and acceptance bands

Boundary integrity

Tier A: R_{obs} per-window ≥ 0.97 , session compliance ≥ 95 per cent; contiguous-green ≥ 120 seconds and ≥ 60 windows; green duty cycle ≥ 90 per cent; longest amber streak ≤ 5 per cent of

session duration; cumulative amber ≤ 7 per cent; per-minute amber cap ≤ 15 per cent; an amber slope alarm drives an automatic tightening ladder with documented reversibility after stability. Precedence across streaks, duty, and amber budgets is declared via a single verdict function.

Tier B: $R_{\text{obs}} \geq 0.95$; compliance ≥ 90 per cent; contiguous-green ≥ 60 seconds; longest amber ≤ 10 per cent; cumulative ≤ 12 per cent; per-minute ≤ 20 per cent.

Tier C: $R_{\text{obs}} \geq 0.93$.

Leakage integrity

Tier A: total variation distance $\leq 1e-6$ with confidence interval upper $\leq 1e-6$; integral floors ≥ 400 events-equivalent per t_{psy} with an uncertainty-quantified crosswalk to counts; count floor ≥ 50 ; true negatives ≥ 200 per class; per-class power and confidence intervals; classes failing floors are nonevidentiary and excluded from aggregates. The true negative protocol is pinned and quality control-audited. Data type sweeps include denormals, not a numbers, flush-to-zero, saturation, mixed-endian payloads, compressed-stream replays, and fused multiply-add differences; failing exemplars are shipped.

Tier B: total variation distance $\leq 5e-6$; integral floors ≥ 250 ; count ≥ 30 ; true negatives ≥ 120 .

Tier C: confidence interval upper $\leq 1e-5$; integral floors ≥ 150 ; count ≥ 20 .

Scale-pyramid monotonicity and anti-alias verification

Monotonicity is defined as sign invariance plus rate-specific total variation distance envelopes; impulse responses, ripple, transition bands, and on-device phase linearity with group-delay variance bounds are verified; any on-device deviation above specification is a refuter. A pass cache, keyed by `rate_filter_hashes`, amortises validation.

Units round-trip and service level agreement

An independent, blind bits $\leftrightarrow$ nats base-swap is executed in a deterministic cold container with randomised order; residual histograms, platform breakdowns, and per-figure unit checksums are produced; the remediation service level agreement is 24 hours with a weekend policy and staged reopen; a red team injects swaps and publishes detection rates; a missed service level agreement demotes to Tier B or pauses claims.

σ^* transport, differential privacy coupling, meta-analysis, hysteresis

Tier A: σ^* coefficient of variation ≤ 7 per cent; two failures demote, two passes reopen, a cooldown is enforced, and a weekly toggle cap is applied. $I^2 > 30$ per cent triggers auto-scoping; $I^2 > 50$ per cent triggers site-scoping; random-effects prediction intervals are included in claim strings; a minimum of ≥ 3 sites is required with a per-site power table. Differential privacy coupling is preregistered as either widened ceilings or increased power; the choice and rationale are recorded in claim strings.

Tier B: σ^* coefficient of variation ≤ 10 per cent; I^2 up to 50 per cent is allowed with narrow badges; a minimum of ≥ 2 sites is required.

Anchors, alias risk, penalties

Anchor alternation cadence M is set with an `alias_risk_simulation`; periodograms with confidence intervals are analysed; cadences within the confidence interval of dominant peaks

are forbidden; a small, preregistered random jitter breaks residual alignment without violating strict_sync or Δt ; pre-lock tests prove no new timing pathologies. Penalty kernels must widen confidence intervals by a minimal fraction; cosmetic penalties fail builds.

Replay_equivalence

Overall replay_equivalence must be ≥ 0.95 and the worst-quantile, typically the decile, must be ≥ 0.93 per metric, with confidence intervals; window segmentation is fixed and hashed; per-window residual histograms and clustering reports expose localised regressions; short-session guardrails declare a minimum N.

Timebase, jitter, slew, daylight saving time and network time protocol drills

Per-device burst classifier thresholds are published; a burst versus drift misclassification audit is performed; jitter budgets and slew caps are set per device class; unannounced, timezone-aware daylight saving time and network time protocol drills are conducted, reporting pass rates; embargo windows prevent immediate accrual after a drill; last-drill timestamps are made visible.

Δt sweeps and indeterminate band

The indeterminate band is defined as $\pm 0.2 \times$ smallest effect size of interest divided by the standard deviation, estimated under autocorrelation with uncertainty; power at the band edges must be ≥ 0.8 ; no invariance language is used when confidence intervals overlap the band; counts and fractions of windows inside the band are printed.

Integrator parity and independence

Integrator parity requires disjoint repositories, compilers, math libraries, random number generators, and step controllers; toolchain fingerprints and link options are included in claim strings; compiler and kernel-binary diffs are provided; parity near χ spikes must be demonstrated on at least two hardware classes; parity claims are blocked without full build metadata.

Weights, exogeneity, caps, leverage

Exogeneity is proven by a weight $\rightarrow$ outcome area under the curve confidence interval that includes 0.5 with an upper bound ≤ 0.54 for Tier A or ≤ 0.56 for Tier B, permutation distributions, multiple lagged shifted-future horizons with a joint null, and an outcome $\rightarrow$ weight regression near zero with a confidence interval; negative controls and lagged placebo predictors are tested. Caps trigger automatic re-power with a visible effective sample size_w delta and a cap-breach-per-hour chart; attempts are limited before mandatory scope-narrowing. Leverage $L_{\max}$ is ≤ 0.08 for Tier A; recurrence thresholds trigger design review and tickets.

Fairness persistence and throttles

Intersectional slices must meet power floors; the $\Delta \wedge_{\text{sync}}$ sign must persist with multiplicity control; chronically underpowered slices trigger declared scope-narrow paths; the decision cadence is throttled while slices are re-powered; throttle events and causes are logged.

Apart baselines, masking placebo, high-dimensional audit

Masking is validated by a Jensen–Shannon distance ≤ 0.03 and a coefficient of variation $_Z \leq 0.10$, plus a high-dimensional classifier audit that fails when residual cues predict conditions beyond a declared bound; a masked-placebo varies ϵ_{macro} while $\Delta\Lambda_{\text{sync}}$ stays within the region of practical equivalence; disagreements between anomaly detection and hash-based checks resolve via a preregistered adjudication ladder.

Matching, overlap, positivity

Balance is assessed by a standardised mean difference ≤ 0.10 or a χ^2 $p \geq 0.10$; weighted effective sample size must be $\geq M_{\text{min}}$ per phase; truncation thresholds are preregistered; tipping-point panels are displayed. After two failed positivity checkpoints, accrual halts; a redesign cookbook, for example with stratification, new recruitment, or instruments, with power impacts and deadlines, is required before resuming.

Smallest effect size of interest, power, alpha-spending

The smallest effect size of interest is ± 0.1 standard deviation for Tier A and ± 0.12 for Tier B; power must be ≥ 0.8 for Tier A and ≥ 0.75 for Tier B, calculated with intraclass correlation coefficient. An O'Brien–Fleming spending schedule is enforced by lints across endpoints with dependency-aware corrections; every look is logged; unscheduled peeks are fail-close for confirmatory outputs and stamp a `look_violation` in the claim string.

ϕ_{attr} governance, canaries, density honesty

Two density classes, such as flow and kernel density estimation, with pre-bounded hyperparameter grids are used; no post-peek expansion is allowed; calibration-only lanes with frozen hyperparameters address drift without peeking; the expected calibration error must be ≤ 0.03 for Tier A, with confidence intervals and bin counts disclosed. The forward–reverse gap $|\Delta\phi_{\text{attr}}|$ must be ≤ 0.03 for Tier A; the remediation order is fixed; maximum attempts are 2. Cumulative sum or sequential probability ratio test canaries have calibrated false alarms, hold-down timers, and freeze only ϕ_{attr} -dependent claims; a dependency graph is printed.

ΔH floors, pooling cap, adaptive windows, detectability, μ deadband

Dual-estimator concordance ϵ must be within $[0.005, 0.01]$ for Tier A; Nemenman–Shafee–Bialek requires ≥ 200 effective counts and Miller–Madow ≥ 100 per window; below-floor windows are nonevidentiary; adaptive window sizes are preferred over pooling; a pooling fraction cap triggers Tier-B demotion; bias corrections and power loss are reported. The minimum detectable ΔQ is computed from sensor noise with a power curve; sub-marginal linkage is treated as sensitivity-only. The μ deadband is derived per session from noise, using median absolute deviation and Gaussian alternatives; per-preprocess and overall sign-flip rates are capped at 2 per cent; a meta-check detects asymmetric protection across sessions.

κ independence and $\kappa \perp$ accountability

The mutual information estimator and small-sample bias correction are declared, for example Gaussian-copula mutual information with permutation calibration; confidence intervals and multiple-testing control apply. If independence is breached, $\kappa \perp$ is introduced with a freeze hash,

and its effect size plus the introduction date are recorded in the claim string; geometry dependence is labelled; a confidence interval-aware rule governs the 0.2 sensitivity fraction threshold; upgrades require revalidation.

β governance, sham, blinding, salience

Powered sham and placebo tests are run per site with blinding checks and operator-salience probes; failure-to-blind rates and retraining logs are published; divergence between human and reinforcement learning results forces domain-scoped badges in titles and claim strings.

χ discipline and clustering

Runs tests, a temporal clustering index with a preregistered critical value, and cap-hit autocorrelation alarms are used; clusters near decision inflections auto-demote to sensitivity and open a masking remediation ticket; an explainability note traces which caps, spans, or modes contributed.

$\hat{E}L$ gating, spoofing, modality isolation

Emotional load gating at 0.75 is justified by receiver operating characteristic analysis and hysteresis; the information_lost metric is linked to $\Delta\Lambda_{\text{sync}}$ bias by simulation; small information loss with significant bias is unacceptable. Synthetic voice and pupil spoofing suites and firmware rollbacks are required; failures pause only the affected modality; modality ablation plots are required before confirmatory rejoin.

ϵ tiers, de-duplication τ , ablation matrix

The de-duplication τ is preregistered per task family with a cross-site stability sweep; post hoc τ edits in confirmatory streams are forbidden. Tier-ablation matrices and τ collision logs are published; dropping any single tier that causes a sign flip is a falsifier; a mechanism audit opens and confirmatory status is demoted.

$\Delta\Lambda_{\text{sync}}$ robustness and fragility

Robust standard errors, either heteroscedasticity-consistent 3 or block-bootstrap, are used; Huber and trimmed estimators are for sensitivity analysis only. A continuous fragility index with a clearly specified contamination model and leave-one-span influence is printed; Tier A is blocked when fragile; remediation to increase robustness is required.

Attribution stability and grouped Owen values

Seed grids and permutation counts are preregistered; convergence curves and an early-stop rule are enforced; seed-to-seed variability is capped; share sums must be within 1 per cent and are linted; instability suppresses mechanism claims and defaults to using Λ_{NS} as the primary estimator.

Λ tracks η over $\|J\|$, local dominance, counter-manipulation

Local elasticities across η deciles are reported, both raw and smoothed; a preregistered dominance metric ensures $\partial\Lambda/\partial\|J\|$ does not dominate in the top $\Delta\Lambda_{\text{sync}}$ deciles. A powered counter-manipulation increases $\|J\|$ at fixed η without rescuing a failed η main effect.

Reciprocity families, priors, spectral leakage

Task families are formally defined; leakage-robust spectral estimators and time-warp replays are linted; segment-length minima and prior sensitivity plots are required; switching priors mid-study requires a fresh preregistration.

Convergence diagnostics and surrogates

Tier A halts on unresolved divergences; rank- $\hat{R}$, energy–Bayesian fraction of missing information, warmup lengths, reparameterisations, and thinning rules are disclosed. Tier B surrogates have a 48-hour offline service level objective for Tier A confirmation; stale provisional badges auto-hide; counts of stale badges are published; lag confidence intervals are checked.

Null–refuter margins and power

Numeric margins for each null are declared with achieved power tables and known-null Type I error calibration; margin breaches fail confirmatory irrespective of signs; unpowered margins are nonevidentiary.

Placebos and auto-tightening without p-hacking

Tightening ladders and thresholds are preregistered; accrual pauses until powered calibration blocks pass under the new thresholds; repeated tighten–peek cycles are forbidden.

Pseudo-dyads, independence, and false positive rate governance

Pseudo-dyad pairs are independent across sites and devices; leakage audits are published; reuse of spans is forbidden. False positive rate bounds have uncertainty quantified; when the upper confidence interval exceeds the bound, auto-tightening triggers; accrual pauses until a calibration block passes.

Instrumental variable suite, discordance, and no swaps

First-stage trajectories are shown; Kleibergen–Paap and Montiel Olea–Pflueger F-statistics, Anderson–Rubin and conditional likelihood ratio intervals, and Hansen's J-test are reported with weak-instrumental variable robust intervals. A predeclared discordance policy for conflicts between J-test rejection and Anderson–Rubin significance is enforced in continuous integration; manual overrides require external co-sign; mid-run instrument swaps are disallowed.

Differential privacy budgets and replays

ϵ -differential privacy, δ , and composition rules are declared; differential privacy queries that access private data consume budget whether or not outputs are published; budget exhaustion hard-stops all confirmatory exports and dashboards; `dp_power_loss` and a precommitted sample ramp with power deltas appear in claim strings.

Missingness missing not at random

A directional robustness claim under missing not at random is required; pattern-mixture and selection-model analyses are run; confirmatory flips force scope-narrow or demotion.

Cross-talk reduction and residual mutual information

Off-diagonal reduction of ≥ 50 per cent with confidence intervals and a residual mutual information bound are enforced using a declared estimator and a secondary method in sensitivity; a tie-break rule is registered.

Adversarial harness costs and latency service level objectives

Per-attack costs and latency distributions, including tail latencies, are published; the Safety-critical profile raises targets and makes failure on a high-harm class a one-strike fail-close; testing cadence is higher for high-harm classes.

Decision aggregator, state machine, and transparency

A machine-readable state log with timestamps, dwell counters, and refuter families is exported; replaying the state machine from these logs must reproduce the dashboard state; mismatches fail continuous integration. nearest_refuter_margin trend alarms notify of creeping fragility; the green state is suppressed during the s+1 dwell period.

Sustainability anchors and penalties

The Anchor_sensitivity_index has a numeric pass or fail threshold; penalties must widen confidence intervals by at least a minimal fraction; any in-bound flip within the declared bounds voids claims, triggers automatic rollback, a cooldown, and a re-anchor calibration plan; external co-sign is required before resuming.

Falsification battery precedence

Any single decisive refuter immediately fail-closes confirmatory claims, irrespective of provisional tags, alpha-spending progress, or dashboard state. A unit test simulates this precedence. Each refuter maps to a nearest_refuter_family; margins are normalised by a standardised σ_r with a published standardisation identifier to ensure cross-site comparability.

Provisional tolerances, sunsets, evidence cadence

At most two provisional tags are allowed per Tier A claim; a retirement plan and deadline to retire at least one tag are required; missed sunsets auto-demote Tier A to B until the review passes. Each provisional item lists target N, sites, sessions, seeds, success criteria, deadlines, and slow-path timings; dashboards show countdowns.

Public red incidents registry and privacy

A public, privacy-reviewed registry publishes minimal structured fields and a linkage to affected claim strings; a takedown policy covers accidental personally identifiable information; entries include a minimal reproducibility snippet to simulate the failure mode; a dwell timer prevents a silent backlog.

Reproducibility on commodity hardware

Wall-clock budgets, memory and graphics processing unit ceilings, and allowed jitter envelopes for scale pyramids and Δt replays are published per hardware_profile_id; mismatching profiles disable determinism claims; reproducibility packs fail when these budgets are exceeded.

Worked end-to-end miniature

A three-site Tier A confirmatory study begins with protocol registration. The team declares a smallest effect size of interest of ± 0.1 standard deviation, an O'Brien–Fleming look schedule with four interim analyses, and an anchor alternation cadence of $M = 20$ windows. The crosswalk_map version 1.3.2 pins all downstream label invariants and power floors, requiring joint power ≥ 0.8 across all items. During preflight, the profile_lock atomically binds the Safety-critical profile, repository commit a1b2c3d, container image digests, kernel modules 5.15.0-91-generic, central processing unit microcode 0xde, central processing unit flags including SSE4.2 and AVX2, basic linear algebra subprograms and linear algebra package OpenBLAS 0.3.23, parallel thread execution and shader assembly compute capability 8.6, compiler flags -O2 -fno-fast-math, random number generator seeds from /dev/urandom with freshness proofs, a dependency graph hash, an environment allowlist restricting to PATH and LD_LIBRARY_PATH, dataset manifests with sampling floors of 100 Hz, feature transforms using a standard scaler version 2.1, anti-alias filters of elliptic order 8, rate_filter_hashes for 0.5 $\times$, 1 $\times$, 2 $\times$, and 4 $\times$ rates, network interface card firmware version 1.89, network time protocol and daylight saving time daemon configs for chrony version 4.3, and entropy sources from a hardware random number generator. Dry-run accrual executes across all three sites simultaneously. The system proves no writer can emit before lock across standard output, standard error, system log, Prometheus metrics, custom JSON exporters, temporary directories, graphics processing unit telemetry via nvidia-smi, buffered caches in Redis, and crash dumpers for core dumps. The dry-run artefact hash e5f6g7h8 is stored. The personally identifiable information redaction policy ensures crash packets contain only the refuter name and manifest hashes. Atomic epoch fencing achieves consensus across 47 devices. The epoch hash i9j0k1l2 is recorded. Fault injection tests prove partial join failures trigger rollback; the rollback manifest m3n4o5p6 restores drivers, compilers, and runtime kernels exactly. The adversarial prose canary runs continuously. Homoglyph triples using Cyrillic a and Latin a, mixed-directionality right-to-left override and left-to-right override tests, code-switch English–Spanish sarcasm, and truncation attacks all pass within the tolerance bands calibrated quarterly. The canary coverage is 98.7 per cent with a false positive rate of 2.1 per cent and a false negative rate of 1.3 per cent. Tolerance bands from the hierarchical calibration model are locked for the release: for Λ , Λ_{NS} , Λ , and Λ_{est} trajectories, the normalised mean square error must be $\leq 1e-4$ and the relative $\Delta \leq 0.1$ per cent; for ϕ_{attr} , the $|\Delta|$ must be $\leq 1e-4$ and the relative $\Delta \leq 0.1$ per cent; for ΔH , the $|\Delta|$ must be $\leq 1e-4$ and the relative $\Delta \leq 0.1$ per cent, with the sign of μ invariant; for χ , the $|\Delta|$ must be $\leq 1e-4$ and the relative $\Delta \leq 0.1$ per cent; for ε tier counts, the absolute Δ must be ≤ 0.5 per cent and the relative $\Delta \leq 1$ per cent. Bands are variance-adaptive for low signal-to-noise ratio strata. Orthogonality quartet tests pass on all latent proxies: permutation $p \geq 0.1$, partial $R^2 \leq 0.01$ with a confidence interval of $[0.003, 0.008]$, conditional mutual information ≤ 0.015 nats with a confidence interval of $[0.010, 0.018]$, and residual correlation ≤ 0.02 . The composite score is 0.89, exceeding the threshold of 0.85. Perturbations, including amplitude scalings of $\{0.5, 1.0, 1.5\} \times$ the standard deviation, phase scrambles, code-switch triads, and homoglyph substitutions, all remain within the realism caps. Worst-case deltas across 15 seeds and amplitude combinations pass the bands. Confirmatory accrual begins. Boundary integrity holds: R_{obs} per-window is ≥ 0.97 across 92 per cent of sessions; session compliance is 96 per

cent; contiguous-green streaks average 145 seconds; the green duty cycle is 92 per cent; the longest amber streak is 4.3 per cent; the cumulative amber budget is 6.1 per cent; the per-minute amber cap is 13 per cent. Amber slope alarms remain quiet. Leakage integrity passes: the total variation distance is $8.7e-7$ with a confidence interval upper bound of $9.2e-7$; integral floors are 420 events-equivalent per t_{psy} ; the count floor is 55; true negatives are 215 per class; per-class false discovery rate and false negative rate are within bounds. Data type sweeps, including tests for denormals, not a numbers, flush-to-zero, saturation, mixed-endian payloads, and compressed streams, all pass. The true negative protocol audit shows 99.1 per cent compliance. Scale-pyramid monotonicity holds at $0.5\times$, $1\times$, $2\times$, and $4\times$ rates. Anti-alias filters meet on-device group-delay variance specifications. The `rate_filter_hashes` validate across all devices. σ^* transport shows a coefficient of variation of 5.9 per cent across sites, well below the 7 per cent ceiling. I^2 is 28 per cent, with prediction intervals of $[-0.12, 0.15]$ standard deviation. The differential privacy coupling choice of "widen ceilings" appears in the claim strings. The nearest refuter margin is 0.34 standard deviations from failure. Orthogonality tests pass with a composite score of 0.91. Perturbations remain within realism caps. Replay_equivalence is 0.975 overall, with a worst-quantile of 0.945. The window segmentation hash `a1b2c3d` validates. The Δt band edges show power of 0.84; 6 per cent of windows lie inside the indeterminate band of ± 0.02 standard deviation. Invariance language is correctly suppressed. Weights pass exogeneity: the weight→outcome area under the curve confidence interval is $[0.48, 0.53]$, which includes 0.5; the permutation p is 0.12; shifted-future horizons show zero lift; the outcome→weight regression is near zero. Caps trigger a modest re-power, with the effective sample size w_{delta} increasing from 1800 to 1950. Leverage L_{max} is 0.065, below the 0.08 threshold. Intersectional fairness slices meet power floors; no sign flips are observed. The masking placebo shows no lift; Λ trajectory drift is within bands. ϕ_{attr} governance uses two density classes, flow and kernel density estimation, with an expected calibration error of 0.019; the canaries are quiet; the forward–reverse gap is 0.011. The dependency graph shows ϕ_{attr} influences $\Delta\Lambda_{\text{sync}}$ but not ΔH . ΔH dual estimators are concordant with $\varepsilon = 0.007$; the pooling fraction is 0.08, below the cap of 0.10; the minimum detectable ΔQ margin is positive at $1.3\times$ the noise floor. The μ deadband yields a 0.8 per cent flip rate across preprocessing variants. κ independence shows a mutual information of 0.014 nats and a partial correlation of 0.18; $\kappa_{\perp}$ is unused. β sham tests are flat with successful blinding rates of 97 per cent. The χ clustering index is 0.12, below the critical value of 0.15. Emotional load spoof suites pass for all modalities. ε -tier ablations show no sign flips when removing any single tier. The fragility index is 0.31 with a minimal contamination fraction of 9.5 per cent. Tier A status is accepted. The compact abstract carries Tier A, the Safety-critical profile, the nearest_refuter_family "boundary", a nearest_refuter_margin of 0.34, I^2 of 28 per cent, a σ^* coefficient of variation of 5.9 per cent, and prediction intervals of $[-0.12, 0.15]$. The claim string validates against schema version 1.3. All figures embed the claim hash. The reproducibility pack builds within commodity hardware budgets: wall-clock time of 4.2 hours, memory of 16 GB, graphics processing unit memory of 8 GB, and jitter envelopes within specification. The continuous integration linter suite returns green across 287 checks. The release proceeds.

Claim string schema fields

The claim string is a comprehensive JavaScript Object Notation structure containing, but not limited to, the following fields: tier, profile, scope_badges, estimator_family, anchor_mode, anchor_cadence, units_roundtrip_flag, sigma_transport_cv, transport_I2, transport_PI, scale_pyramid, dt_parity, integrator_parity, weights_auc_CI, outcome_to_weight_regression_CI, leverage_max, fairness_delta, minority_retention, intersectional_retention, baseline_status, power, alpha_spending, phi_forward_reverse_gap, phi_ece, phi_canary_status, deltaH_concordance, pooled_windows_count, pooled_fraction, mu_deadband_stability, kappa_MI, kappa_partial_corr, kappa_perp_sensitivity, beta_safety, chi_caps, chi_runs_flag, chi_clustering_index, el_info_loss, epsilon_sum, epsilon_tier_drop_sign_stability, deltaLambda_fragility_index, attribution_share, attribution_stability, lambda_tracks_eta, reciprocity_bf, reciprocity_priors, nulls_pass_margins, placebo_rope, pseudo_dyad_fpr, iv_stats, dp_budget, dp_power_loss, transport_meta, replay_eq_overall, replay_eq_worst_decile, anchor_sensitivity_index, gradient_priors_card_id, agnostic_audit_band, agnostic_audit_power, label_shadow_eval_id, label_swap_delta_id, crosswalk_schema_version, crosswalk_dependency_graph_id, crosswalk_joint_power, window_remap_hash, egress_audit_id, atomic_epoch_hash, dry_run_artifact_hash, pii_redaction_spec_id, monotonic_counter_registry_id, independence_score_fn_id, independence_score_shap_id, independence_supply_chain_penalty, degraded_budget_quarter, sign_cadence_model_id, sign_anomaly_fpr_fnr, hidden_seed_pack_id, hidden_seed_escrow_attest, dataflow_coverage, mutation_catalog_id, placebo_grid_spec, rollback_manifest_hash, mi_audit_ceiling, nondet_disallow_list_id, band_hier_model_id, band_change_point_id, adversarial_canary_cov, orthog_composite_score, realism_cap_spec_id, manifest_ladder_rules_id, crash_tombstone_id, attestation_freshness_max_age, spoof_challenge_id, quarantine_latency_slo_ms, thermal_bias_link_fn, recalibrating_label_dwell, epsilon_caps_release_quarter, overlap_bootstrap_plan, independence_min_margin, schema_extension_allowlist_id, abstract_compact_id, dashboard_claim_hash, boundary_verdict_fn_id, amber_tightening_ladder_id, tn_protocol_id, dtype_corner_matrix_id, anti_alias_device_verif_id, unit_swap_redteam_rate, sigma_dp_choice, alias_risk_periodogram_id, alternation_jitter_budget, penalty_min_widening, segmentation_hash, burst_classifier_id, drill_pass_rate, dt_band_power_curve_id, toolchain_fingerprint_id, kernel_binary_diff_id, exog_joint_null_id, cap_breach_rate_chart_id, leverage_recurrence_threshold, fairness_scope_narrow_recipe_id, high_dim_mask_audit_id, positivity_redesign_cookbook_id, of_dependency_correction_id, phi_canary_hold_down, density_calibration_lane_id, adaptive_window_spec_id, dQ_detectability_margin, mu_deadband_method, kappa_mi_estimator_id, kappa_perp_freeze_hash, beta_blinding_audit_id, chi_cluster_crit_value, el_spoof_suite_id, tau_stability_sweep_id, fragility_contam_model_id, attribution_convergence_curve_id, eta_elasticity_panels_id, reciprocity_prior_sensitivity_id, surrogate_stale_count, null_margin_power_table_id, tightening_ladder_log_id, pseudo_dyad_leakage_audit_id, iv_discordance_policy_enforced, dp_budget_state, dp_power_re_ramp_id, mnar_directional_result, crosstalk_secondary_estimator_id, adversarial_latency_slo_id, state_machine_log_hash, anchor_penalty_widening, falsifier_precedence_test_id, provisional_retire_plan_id, red_incident_privacy_policy_id, commodity_budget_profile_id, jitter_envelope_id.

This completes the Λ - Ψ Indifference Protocol (PIP) v1.3 specification.

The maths does not belong to me, but to everyone.

Λ - Ψ Indifference Protocol (PIP) v1.3 by Jordan Lee McDonald is licensed under Creative Commons Attribution 4.0 International.

Λ - Ψ Harmonic Reciprocity Protocol v1.1

Tier 0: Stance, Leakage, WORM, and Dry Checks

P_obs only stance is enforced under PIP v1.3. The indifference principle is applied. The pip_version, pip_lock_hash, and completed pip_checklist with timestamps and operator signature hash are written to WORM. The machine-readable schema_hash is included.

Leakage detection uses planted events with a preregistered schedule. The minimum and maximum inter-event intervals, size distribution family, and seed derivation path from the root seed are published. Operator overrides during mains are forbidden. Precision and recall must have 95% Wilson CI lower bounds ≥ 0.99 and ≥ 0.98 respectively. The exact CI bounds and sample size n for every sentinel are printed. A per-epoch plot of precision and recall with control limits is generated.

The total_ref_count is unique-by-hash per window across the entire run. The raw_count and duplication_ratio = raw_count divided by unique_count are reported per window; alerts trigger if duplication_ratio > 1.2. Duplicate hashes are rejected; the first-seen path and metadata are logged for provenance debugging. NTFS alternate data streams and resource forks are detected.

The exact hash function is SHA-256. Path canonicalisation uses NFKC normalisation, case folding, symlink resolution, path separator normalisation across OSes, and accounts for locale and filesystem case-sensitivity. Tests for UNC paths, junctions, bind-mounts, IPv6, and DoH/DoT are included. A collision test suite and a list of normalised path pairs that collapse are published; warnings are issued for any non-ASCII that normalises. The path_canonicalizer_hash is stored.

The allowlist has a time-to-live of 300 seconds. The cache_age and next-refresh ETA are printed. A soft warning triggers at 80% TTL, initiating a background refresh task whose success or failure is recorded. A hard fail occurs at TTL+. The last refresh cause, manual or auto, is printed. The allowlist_cache_hash is included. The ASN allowlist is snapshotted and hashed daily; changes mid-run without a signed change record cause failure.

Network sentinels detect sub-kilobyte leak probes, IPv6, QUIC/HTTP3, DNS-over-HTTPS/TLS, and localhost tunneling. Protocol, host, port, bytes, ASN, SNI, and certificate hashes are

recorded. Frequency-gating prevents floods; if the rate exceeds a threshold, the system auto-pauses and requires operator acknowledgement. Only preregistered ASNs and ports are allowed. Any unauthorised egress fails the run.

Dry-run boundary batteries execute every 30 minutes during mains and after environment changes. Pass stamps are logged.

WORM uses immutable, append-only storage with a backend configuration hash and a signed attestation of object-lock, retention mode (governance or legal hold), and legal hold features. The object-lock expiration is printed. Merkle proof tree height, leaf-order hash, and an erasure-coding checksum if used by the backend are recorded. Replay equivalence is defined as ordered event-hash equality with a hashed policy for normalisation transforms, e.g., timezone unification, float formatting. Failing fields must be zero outside this whitelist. Per-field mismatch rates and first differing indices are printed.

The daily canary runs between 00:00 and 02:00 UTC with a backup window until 04:00 and a jitter tolerance. The expected failure signature hash and `canary_last_seen` are stored. The observed failure signature hash must match the expected digest. The canary's full stderr and stdout digest are stored. Missed-window remediation is documented.

Monotonic clocks timestamp all leakage events; wall-clock offsets are recorded. Absolute drift must be ≤ 100 ms, drift rate ≤ 1 ms per minute, and max jump per event ≤ 10 ms. The clock source and stratum are recorded.

Tier 1: Metric, Transforms, Invariance, and Lint

The whitening map is $s' = W(s - \mu)$. In whitened space, M_w is I . In original space, $M = W^T W$. The condition number $\text{cond}(W)$ must be $\leq 1e3$ and $\sigma_{\min}(W) \geq 1e-6$. The expected bound $\text{cond}(M) \approx \text{cond}(W)^2$ is derived, printed, and gated explicitly; inconsistencies are flagged. M must be symmetric positive definite; a Cholesky check with jitter fallback is added. The smallest eigenvalue and fraction of near-zero eigenvalues are printed. Ledoit-Wolf shrinkage is applied if $N/d < 10$; its randomisation and α shrinkage grid are frozen and hashed. The SVD tolerance is $\max(\epsilon \cdot \max(m,n) \cdot \sigma_{\max}, 1e-12)$. The relative Frobenius norm $\|M - W^T W\|_F / \|W^T W\|_F$ must be $\leq 1e-10$, with a machine epsilon multiplier per platform.

Transform rules: $g_E' = W^{-T} g_E$, $v' = W v$, $\nabla \Psi' = W^{-T} \nabla \Psi$. Λ_{mesh} and η invariance are proven. The `metric_choice` flag selects `euclidean_whitened` or `metric_dot`. Both pipelines run on validation batches. Metric equivalence requires slope in $[0.98, 1.02]$, intercept CI includes 0, $R^2 \geq 0.99$, and concordance correlation coefficient for both Λ_{mesh} and $\eta \geq 0.995$. Bland-Altman limits of agreement must have mean bias near zero and LoA within $\pm 0.5\%$ of the mean or within $3 \times \text{SE}$ of the difference, expressed as a percentage of dynamic range.

Invariance demos run per window with stamped RNG, direction sampler, and a sample count automatically adjusted to achieve ≥ 0.9 power to detect a 0.2% bias. $|g_E \cdot v - g_E' \cdot v'|$ and

relative error use tolerance $\max(5e-4, 3 \times SE)$. Two consecutive breaches halt the run. The linter checks axis labels, figure metadata, PDF embedded text, SVG text nodes, and performs OCR on rasterised figures to catch unit strings in images, printing confidence and suspected positions. Only "nat·s⁻¹" and "J·s⁻¹" are allowed. Banned phrases include "per nat" and variants. Unicode homoglyphs are scanned with NFC/NFKC pre-normalisation and collision tests; offending code points and their names are printed with a diff snippet. The `gradient_domain_demo_hash` is published, including a strict non-orthogonal W example with numerical and symbolic invariance verification.

Tier 2: Variables, Units, Signs, Chain Rule, and Micro-Checks

Ψ is dimensionless. $\nabla \Psi$ has units state^{-1} in whitened coordinates. The linter blocks "per nat" and "nat⁻¹". The g_E sign convention uses `ge_sign` in $\{+1, -1\}$. A synthetic flip test verifies Λ_{mesh} sign inversion. Channelwise sign checks are performed with planted flips of defined prevalence and effect sizes; ROC requires $AUC \geq 0.95$, sensitivity ≥ 0.9 at 5% FPR, and power ≥ 0.9 to detect flips. Per-channel ROC curves are printed. Automated localisation reports channel ID, magnitude, and Λ_{mesh} impact.

If $E = E(s, t)$, $g_E(s, t) = \partial E / \partial s$ at fixed t . The magnitude distribution of $\partial E / \partial t$ and the fraction of windows where $|\partial E / \partial t| > 0.1 |g_E \cdot v|$ are published. A temporal-energy share metric is computed and trends are shown across conditions. If the fraction exceeds 10% in a block, the block is auto-demoted to sensitivity.

The chain-rule check: $dE/dt - g_E \cdot v - \partial E / \partial t$. Residuals use tolerance $\max(5e-4, 3 \times SE, 5e-4 \cdot IQR(dE/dt))$. Block bootstrap CI with fixed block length of $5 \times$ the largest autocorrelation lag is used; sensitivity to $3 \times$ and $7 \times$ is printed; an ACF plot and chosen lag are provided. Huber residuals are used in the CI. Two consecutive breaches fail.

$v = ds/dt$ in whitened space. The differentiator (central) and smoother kernel are fixed; an L-curve or GCV plot for smoothing is provided; alternatives are sensitivity only; decision invariance under those alternatives is printed. Δt sweeps $\{0.5, 1, 2, 4\}$ show η invariance bands. Overlap is quantified, effective sample size is computed, and the CI method is adjusted accordingly. J_{vec} is auxiliary and blocked from Λ_{mesh} .

Zero-norm masking: if $\|v\| < \text{eps}_v$ or $\|\nabla \Psi\| < \text{eps}_{\text{grad}}$, where eps_v and eps_{grad} are $1e-6 \times$ the per-run median norms, scaled per-dimension by eigenvalues of M_w to account for anisotropy, then η and ϕ_{attr} are masked. Mask rates and overlap with low Λ_{mesh} windows are reported.

Finite-difference uses random unit axes and ϵ in $\{1e-4, 5e-4, 1e-3, 5e-3\}$ scaled by per-axis SD. Monotone convergence to slope 1 is required. A second-order check ensures the first-order regime holds; flags are raised if curvature dominates at smallest ϵ . The slope vs ϵ log-log plot must have residual slope in $[0.95, 1.05]$. Plateaus with slope change $< 1e-3$ flag potential quantisation. Automatic differentiation on a proxy model requires $\geq 95\%$ directional sign

agreement and median relative magnitude error $\leq 5\%$; top-10 directional disagreements with axis id, ϵ , and local norm are listed.

Tier 3: Hypotheses and Phrasing

- H1: Cadence peaks at f_{res} with Δ_{harm} within ϵ_{eq} .
- H2: $\sigma \propto g_E \cdot v (k_B T_{\text{eff}})^{-1}$ with zero intercept.
- H3: σ positive on average; v_{loop} tracks confusion.
- H4: $1/f_{\text{res}}$ drifts toward $1/T_k$ monotonically.
- H5: η outpredicts single v channels; σ tracks η .
- H6: $\partial\sigma/\partial\|v\| > 0$ in math-on windows.
- H7: $R_{\text{obs}} \geq 0.97$ and anti-circularity.
- H8: Bits $\leftrightarrow$ nats and Λ_{mesh} to σ round-trips pass.
- H9: Decisions invariant under θ_{cos} and θ_{ψ} .
- H10: σ^* stable across modalities.
- H11: Conclusions unchanged under σ_{min} perturbations.
- H12: γ_{σ} provisional; no primary depends on it.

The hypothesis registry is locked with hypothesis_gates_hash. Priors and seeds are published.

Tier 4: Formal Analysis Plan by Hypothesis

H1: Generalised means M_r compared with priors. The exact functional form of the prior on r and an alternative are published; mid-run tuning is forbidden. Prior predictive overlays are printed. Bridge sampling for Bayes factors. Minimum log-evidence ESS ≥ 1000 and MCSE ≤ 0.5 ; otherwise auto-switch to stepping-stone and record the temperature ladder and $\Delta \log Z$. Δ_{harm} with bootstrap CI. Gate: $\text{BF}(r=-1) \geq 10$ and Δ_{harm} within ϵ_{eq} . The CI method for SE_boot and its seed are locked. Every Δ_{harm} breach maps to a detector class: aliasing, latency drift, line noise, VRR, jitter misconfig; unknown is disallowed. Detector scores and thresholds are attached. Per-condition ϵ_{eq} is printed with CI method and seed; near-threshold windows are highlighted.

H2: Huber regression of σ on $g_E \cdot v (k_B T_{\text{eff}})^{-1}$. δ is preregistered with a numeric value and MAD-based derivation; slope stability across $\delta \in [0.5, 2] \times \text{prereg value}$ is reported; alternates are sensitivity. Errors-in-variables sensitivity via Deming or SIMEX is considered. Slope and intercept invariance across δ band. Strata: task, modality, participant. Heterogeneity with Cochran's Q. Sandwich SEs for heteroskedasticity. Gate: intercept CI includes 0, $\text{BF} \geq 3$ for zero-intercept, and posterior density at zero ≥ 0.25 of the mode. Leave-one-channel-out: slope change $\leq 20\%$; catastrophic changes $> 50\%$ or CI crossing zero auto-demote the channel. Spline test for nonlinearity; if significant, confirmatory proportionality claims are withdrawn and documented.

H3: One-sample t-test and Wilcoxon signed-rank test. Variance-stabilising transform if skewed. Loop detector: Savitzky-Golay (11,3; fallback 7,2). Deadband $\epsilon_{\sigma} = q \cdot \text{MAD}$, q in $\{2, 2.5, 3, 3.5\}$.

Logistic link with preregistered covariates and weakly informative priors, e.g., Normal(0, 2.5) on logit scale. Calibration: Brier score and ECE with bootstrapped CIs; reliability curves with isotonic recalibration overlay. Planted synthetic loops across a specified SNR grid with a defined injection mechanism; TPR ≥ 0.9 at FPR 0.05 with Wilson intervals is gated.

H4: Kendall's τ and isotonic regression. Dwell duration and T_k steps preregistered. BY correction across T_k and H1-H6 families; adjusted p-values are presented alongside BFs. Power under observed jitter and n is computed; achieved power per block for τ is printed; underpowered blocks are automatically demoted. Coherence pre-gate: only test τ when coherence $\geq$ threshold; skipped cases are logged; coherence gate pass rates are included. Gate: $BF \geq 10$ or $p \leq 0.01$ with $\tau \geq 0.3$.

H5: Nested ΔR^2 or ΔAUC with subject- and time-blocked CV. Permutation count ≥ 5000 . Hold-out stability ≤ 0.05 . Leakage audit hashes group IDs in CV splits and prints any overlap; fails on leakage. The effect on $\Delta R^2/AUC$ if leakage is hypothetically introduced is shown. Control: shuffled $\nabla \Psi_h$ must cause $\geq 90\%$ drop in ΔR^2 or AUC; otherwise alignment claims are suspended.

H6: Isotonic slope of σ vs $\|v\|$ across Δt sweeps. Gates: $\theta_{\max}$ in $\{45^\circ, 60^\circ, 75^\circ\}$ or $\eta > 0.5$. Per-gate sample sizes and achieved power are disclosed; power < 0.8 is sensitivity. Partial regression controlling for T_{eff} ; σ residuals after partialing out T_{eff} are plotted; monotone slope and CI on residuals are computed.

H7: R_{obs} as held-out log predictive density ratio. The exact scoring rule, held-out fraction, and per-sample normalisation are printed and fixed across sites; deviations are reported. CI with BCa or block bootstrap. Ladder: $R_{\text{obs}} \geq 0.97$ confirmatory, 0.95-0.97 sensitivity, below 0.95 exclusion. Anti-circularity: symmetry index $|SI| \leq \tau_{\text{sym}}$; SI is defined mathematically; null distribution via label-swapping; τ_{sym} and CI are preregistered.

H8: Control chart for σ_{bits} divided by $\sigma_{\text{nats}} \ln 2^{-1}$. Western Electric rules: 1 beyond 3σ , 2 of 3 beyond 2σ , 4 of 5 beyond 1σ , 8 on one side. Rule counts and false-alarm rate under block dependence via phase randomisation are annotated on control charts. Raw triplets $[\wedge_{\text{mesh}}, k_B T_{\text{eff}}, \sigma]$ saved.

H9: Bitwise assertion for H1-H6 pass/fail flags. `base_swap_report` with empty diff. Both run hashes and a visual diff panel are stored.

H10: r_m as ICC(3,1) with two-way mixed effects, absolute agreement. CIs and within-subject SD are reported. Leave-one-out divergence. Justify session spacing.

H11: σ_{min} perturbations $\times \{0.1, 1, 10, 100\}$. TOST equivalence with preregistered numeric margins and α .

H12: γ_σ priors log-uniform [1e18, 1e22]. Prior predictive checks. Tracks A, B, C. Dependency booleans no_primary_depends_on_gamma_sigma and no_primary_depends_on_sigma_min must be true, proven by static DAG analysis.

Tier 5: Harmonic Operator, Domains, Guards, and Controls

$f_{\text{res}} = 2 f_1 f_2 (f_1 + f_2)^{-1}$. $T_{\text{res}} = 2 T_1 T_2 (T_1 + T_2)^{-1}$. $\delta_T = 4 T_1 T_2 (T_1 + T_2)^{-2}$. $\phi_f = 4 f_1 f_2 (f_1 + f_2)^{-2}$. $\epsilon_{\text{eq}} = \max(1e-12, 5e-4, 3 \times \text{SE}_{\text{boot}})$.

Residuals: Δ_{harm} per condition with bootstrap CI. $T_1 = T_2$ lines on all plots per participant; Δ_{harm} at machine precision with CI; fail if absent or inflated.

Domain guards: T_{min} , T_{max} , f_{min} , f_{max} from device specs with hashed provenance. Clipped-trial counts printed. Lower bounds enforced to avoid division by zero.

Reciprocity dwell: Minimum duration and T_k steps preregistered.

Tier 6: Sampling, Spectra, Aliasing, Jitter, and Latency

Sampling rate $\geq 10 \times \max f_i$. Anti-alias LPF below 0.45 Nyquist with certified specs. Multitaper: NW and K stamped.

Spectral leakage: 1/f detrending. Rayleigh resolution $1/T_{\text{record}}$ for low-f exclusion; the computed exclusion width per block is printed and enforced; a minimum width is applied.

Hardware lines: Exclude 50/60 Hz, refresh rates, harmonics. VRR detection from photodiode/GPU timestamps; refresh histograms and GPU timestamp traces are recorded; if VRR is active, it is either disabled or exclusion bands are widened.

Jitter families: Uniform primary; Poisson, Gaussian, lognormal sensitivity. Effect-size drift of δ_T and ϕ_f . A decision table with power vs effect-size drift and seeds is published; the chosen family and level are stamped with rationale.

Latency governance: Auto-pause on drift > 1 ms/minute. PTP/NTP offset and jitter per device are recorded; auto-pause if sync error > 0.5 ms. Resync cap per block, e.g., ≤ 2 ; exceedance demotes. Pre/post latency distributions: median, IQR, 95th. Drift trends are plotted.

Tier 7: Preprocessing, Channels, and S_{EL}

Preprocessing hashes for pupil, EDA, HRV, EEG. Mismatch fails build. Parameter diffs are published on mismatch.

Pupil: Luminance calibration, blink interpolation, saccade regression. Lag compensation; residual lag > 10 ms flags sensitivity.

EDA: Tonic/phasic separation, temperature correction.

HRV: RR preprocessing, ectopic handling, minimum beats per window.

EEG: Notch filter, bands, rereference, ICA, artifact rejection.

S_EL independence: $I(S_EL; \sigma) \leq 0.02$ nats or $|\text{partial } r| \leq 0.2$. Kraskov k and bias correction are specified; $n/k \geq 20$ required; MI CI via block bootstrap is reported and gated. If $MI > 0.05$ nats, $S_EL \perp$ enforced.

Tier 8: Clocks, Decorrelation, α_i , and Demotion

Decorrelation: Partial correlations or PCA. Pre/post correlation matrices with CIs. Residual $|\rho| \leq 0.2$. PCA fallback.

α_i estimator: Kraskov kNN with stamped k . Permutation CI; width ≤ 0.15 and n/k sufficiency required; otherwise ridge fallback with λ chosen via cross-validated hash; sensitivity label applied. Hierarchical α_i with Dirichlet 1 prior.

Ridge fallback if α_i CI width > 0.15 or $TVD > 0.2$. Non-harmonic after two reciprocity failures or α_i instability. Recovery requires clean pass and minimum dwell; dwell is logged.

Tier 9: Ω Integrity, Guard Region G , and Anti-Circularity

R_{obs} : Held-out log predictive density ratio. CI with BCa or block bootstrap. Ladder: $R_{\text{obs}} \geq 0.97$ confirmatory, 0.95-0.97 sensitivity, below 0.95 exclusion.

Guard region G : Convex-hull or density-quantile choice is preregistered. Coverage printed. σ and $\nabla \Psi_h$ effect deltas with and without G plus CIs are reported; stability is gated.

Anti-circularity: γ_h frozen. Forward and reverse scores, reverse KL, symmetry index. $JS \leq 0.03$ nats, $CV_Z \leq 0.10$.

Tier 10: T_{eff} Provenance, β Contraction, and Safety

Constants: k_B and $\ln 2$ locked with constants_hash.

T_{eff} : Landauer scaling with sensor model, anatomical site, ambient correction. CI propagated via delta method to σ and ρ_P , ρ_ψ bands. An independent calibration, e.g., calorimetry or manufacturer calibration, is used; agreement within an a priori margin is gated.

Contraction law: ΔH vs $\beta T_{\text{eff}}^{-1}$. Expected slope $-1 T_{\text{eff}}^{-1}$. Overlay ± 1 SD bands on ΔH slopes. Participant-level slope histograms and mixed-effects slope and variance components

are reported. Safety: $\partial H_{\text{neu}}/\partial \beta < 0$, $\partial RT/\partial \beta < 0$ with mixed-effects. Effect sizes are reported. $\leq 5\%$ adverse changes in $\hat{C}_o$ and $\hat{E}L$. Any protocol pauses or adverse events are published with WORM-stamped incident reports.

Tier 11: MCMC, PSIS, Numerics, and Precision

Whitening QA: Axis RMSE ≤ 0.05 SD, angle drift $\leq 10^\circ$, $\text{cond}(W) \leq 1e3$, $\sigma_{\text{min}}(W) \geq 1e-6$. Ledoit-Wolf if $N/d < 10$.

Density and score: Adaptive kNN, spline score matching, frozen λ grid.

MCMC: ≥ 4 chains, independent seeds, rank plots, acceptance 0.65-0.8, bulk/tail ESS ≥ 2000 , split $\hat{R} \leq 1.05$. Divergences trigger auto-rerun with documented parameter changes. Posterior drift is checked via Wasserstein distance or rank plots; promotion only if within prereg thresholds.

PSIS: Pareto $k < 0.7$. Tempering auto-triggers; level, tempered ESS, bias vs baseline are logged. Hill tail-index is computed and gated.

Numeric stability: Float64 enforced at module boundaries and deep within kernels. FMA behaviour is pinned via compiler flags; CPU features are recorded. Underflow/overflow totals and largest magnitude per block are printed. Kahan/Neumaier summation for $N > 1e5$. Dtype audit table is published.

Units round trip: $[\Lambda_{\text{mesh}}, k_B T_{\text{eff}}, \sigma]$ triplets saved for failures with a verification script and verification_hash. Two consecutive breaches halt. Bits $\leftrightarrow$ nats conversions are verified with explicit base markers; a unit test fails on log base mismatches.

Tier 12: Naming, Styles, and Banned Phrases

Canonical identifier: sigma_rate. Linter blocks variants via AST analysis and confusable Unicode in identifiers; offending tokens are printed. UTF normalization NFKC. Homoglyph scanning.

Tier 13: Cadence Coherence and Proximity Gates

Coherence threshold > 0.3 , tied to a permutation-based null distribution; empirical α is published. Half-bin proximity rule: threshold is $\min(\text{half Rayleigh}, 0.02 \text{ Hz})$ explicitly; justified and locked. Spectral annotations: Rayleigh resolution. F-test with multiple-test correction.

Tier 14: Clocks, Bases, and Planck Context

Base invariance: H1-H6 rerun under θ_{cos} and θ_{ψ} . base_swap_report with empty diff. Both run hashes stored.

CODATA: Year and t_P stamped. Uncertainty propagated to ρ_P bands.

Planck context: ρ_P and ρ_ψ with "context only" watermark. Static analysis proves they are never used as features or selectors; a signed proof artifact is stored.

Tier 15: Exclusions, Thresholds, Bias, and Randomisation

Participants: Exclude if >20% technical failures, >30% nonharmonic blocks, >30% R_{obs} failures. $\pm 5\%$ sensitivity analysis. Robustness badge as Jaccard similarity of pass/fail vectors; stability ≥ 0.9 required. Inverse probability weighting and doubly robust sensitivity for missing not at random exclusions are run; divergence is reported in the executive table.

Blocks: Exclude if $R_{\text{obs}} < 0.95$.

Windows: Nonconfirmatory only.

Exclusion ledger: Hashed IDs and reason codes. Reason code ontology is versioned; 100% mapping is enforced; remediation hints are printed; an audit table is included.

Randomisation: Latin or Williams design with hashed matrices. Balance tests are run. First-exposure blocks removed if order impact. The impact of any reweighting on primary estimates is shown.

Tier 16: Mapping Tables, Uncertainty, and Transport

Sensor to latent mapping: Calibration constants, sampling, lags.

Site transport: Δ_{harm} and $\Lambda\text{-}\eta$ coupling CI overlap. TOST equivalence with preregistered bounds. Sites failing either are demoted; exact violations are printed.

Tier 17: Reciprocity Demotion and Recovery

α_i instability: CI width > 0.15 or TVD > 0.2. Non-harmonic after two failures. Recovery: clean pass and minimum dwell.

Tier 18: Jitter and Negative Controls

Jitter: Most conservative by power, chosen from a decision table with seeds.

Negative controls: $T_1 = T_2$ lines on all plots per participant. Δ_{harm} at machine precision with CI; fail if absent or inflated.

Tier 19: γ_σ Governance, Priors, Tracks, and Booleans

Priors: log-uniform $[1e18, 1e22]$. Prior predictive checks.

Track A: Hierarchical homogeneity, $I^2 \leq 30\%$ with CIs.

Track B: Held-out families; calibration plots and coverage.

Track C: κ_c and $g(\beta, T_{\text{eff}})$; residual curvature independence.

Dependency booleans: `no_primary_depends_on_gamma_sigma` and `no_primary_depends_on_sigma_min` true, proven by static DAG.

Tier 20: Executive Outputs, Operator Checklist, Tolerances, and Release

Executive table: Pass/fail per hypothesis, effect sizes, BF, p, CIs, reason codes.

Operator checklist: R_{obs} , JS, CV_Z, ESS, $\hat{R}$, PSIS k, units round-trip, Δ_{harm} . Near-threshold passes require timestamped, signed acknowledgements; absence blocks publication.

Tolerance schema: JSON with ϵ_{eq} , round-trip tolerances, etc. Versioned changelog. Regression tests must pass against prior schema hashes; diffs are printed.

Synthetic sanity suite: $T_1 = T_2$, $T_2/T_1 \rightarrow 0$, $T_2/T_1 \rightarrow \infty$, T_2/T_1 in $\{0.25, 4\}$. Bisect-with-seeds tool. Minimal repro packets stored.

Units_inference_table: Symbols, units, conversions, pass/fail. No NA placeholders. A fail-fast gate triggers if any symbol lacks a unit or conversion.

Figure provenance: `manifest_id`, `figure_hash`, `tolerance_schema_hash`, `preprocessing_hash`. Cross-checked in CI; mismatches block publication; a concise diff of embedded metadata is shown.

Privacy: PHI-scrub with regex, structured detectors, and multilingual, context-aware patterns. FP/FN rates with CIs on a held-out red-team corpus are published; de-identification residual risk is shown.

Replay: `replay.json` with ≥ 0.95 equivalence, environment lockfile. A single-command CI recomputation bundle is required to recompute H1–H6, Ω ladder, and contraction gates; fail if gate disagreement $> 5\%$ or any figure/table hash mismatch.

Data release: Includes SESOI calculators, bits $\leftrightarrow$ nats toggles, Ω ladder, contraction gates. Replay ≥ 0.95 required; command and hashes published.

Negative-Episode Telemetry and Policy

Onset-to-policy latency, recovery proportion, 90th-percentile recovery time with CI, inter-episode interval. Kaplan-Meier curves with confidence bands. Censorship reasons. Stratification by β schedule.

Visual Discipline and Units

Axes: Λ_{mesh} in $\text{J}\cdot\text{s}^{-1}$, σ in $\text{nats}\cdot\text{s}^{-1}$. Δ_{harm} equality bands. v_{loop} with confusion EMA. ρ_P and ρ_{ψ} with uncertainty.

Final Stance

Operational $\Lambda = \sigma$ in $\text{nats}\cdot\text{s}^{-1}$. $\Lambda_{\text{mesh}} = g_E \cdot v$ invariant. Harmonicity and reciprocity by BF and Δ_{harm} . Alignment and monotonicity with leakage guards. Integrity and anti-circularity machine-checked. Numeric precision enforced. Planck bands context only. γ_{σ} provisional. All tolerances and constants PIP-locked, WORM-anchored, replayable. Seed governance uses a single root seed stored in WORM; a derived seed map for modules is printed; rotation is via a locked procedure. Environment drift is pinned; compiler flags, BLAS threads, deterministic kernels are fixed and compared at replay; drifts cause failure and diffs are printed. Adverse event safety is logged. Jitter family choice is stamped. Documentation hashes $\text{gradient_domain_demo_hash}$, $\text{base_swap_report_hash}$, $\text{leakage_detector_suite_hash}$ are stored. Power transparency: achieved power per hypothesis is published; confirmatory claims are gated at ≥ 0.8 ; underpowered results are demoted.

The maths does not belong to me, but to everyone.

© 2025 Jordan Lee McDonald. Λ – Ψ Harmonic Reciprocity Protocol v1.1 is licensed under the Creative Commons Attribution 4.0 International License (CC BY 4.0).