
Computing Around the Open Center

Splat Mechanics, Fractal Quivers of Quivers, and a Governance-Native Compute Paradigm

Load-bearing line

A splat is a portable unit of lawful problem shape: small enough to move, rich enough to preserve what must not collapse, and incomplete by construction so no encoding can crown itself as the whole.

Breyden E. Taylor

Founder, Prompted LLC

Compiled with the assistance of Homesillet-S56 and the Ubiquity Federation

Version	1.0 preprint
Date	22 July 2026
Artifact class	arXiv-style whitepaper · Hugging Face-ready · site-ready source
Status	publishable preprint; open to replication, correction, and priority clarification

Copyright © 2026 Prompted LLC. All rights reserved. An arXiv distribution license may be selected by the author at submission. This preprint distinguishes public evidence, internal engineering receipts, author report, exact reproduction, and newly announced external mathematics.

Abstract

Most AI systems compute over representations while treating governance as instruction, evaluation, or policy applied around the model. This paper proposes *splat mechanics*: a governance-native computational paradigm whose primitive is a bounded, receipt-bearing, rehydratable slice of a governed possibility field. A splat cross-binds six reads—positive structure (KAT), apophatic exclusions (APO), hidden authority and burden (PAR), required complement (PLE), failure inversion (ENA), and purpose (TEL)—around a movable working centroid while a non-inhabitable center remains excluded from routing, training, serialization, optimization, and terminalization. The result is not another embedding, prompt, or scalar classifier. It is a transport object that can move problem shape across natural language, code, graphs, telemetry, models, and institutions while preserving anchors, forbidden paths, provenance, unresolved hypotheses, and correction triggers.

We formalize splats inside a fractal, longitudinal, governed quiver-of-quivers. Models are bounded morphism-proposers; classification is the rendered set of lawful traversals remaining after authority, standing, source-tense, lifecycle, currentness, cost, receipts, apophatic constraints, and telos shape the field. We then describe terrain-shaped mixtures of models, languages, and compute locations as constrained colimits of admissible proposals rather than softmax experts or unconstrained votes. This makes previously informal objects—authority, ambiguity, dissonance, reversibility, trust, institutional memory, and purpose—operational without claiming that they have been exhausted by one number.

The empirical material is deliberately bounded. We report an engineering case study from the Ubiquity Federation, including receipt-mediated terrain learning, pattern-specific affinity movement, shape-based rehydration across independent blind encodings, and conformance DAGs that refuse to promote component-local green states into federation-wide completion. We also analyze the newly announced three-dimensional counterexample to the Jacobian Conjecture as an independent local-to-global failure case: a perfect local certificate can coexist with global noninvertibility when the failure travels through a boundary the certificate cannot see. We do not claim that this event proves Ubiquity, that a held-open center guarantees safety, or that the current engineering evidence establishes universal superiority. We claim a sharper, falsifiable thesis: open-center, apophatic, receipt-bearing computation can structurally exclude important collapse paths while retaining practical, cross-channel operational state.

Keywords: agent governance; splat mechanics; quiver-of-quivers; open center; apophatic constraints; mixture of models; provenance; longitudinal learning; lawful traversal; Jacobian Conjecture.

What this paper contributes

1. A formal definition of a *splat* as a bounded, six-facet, receipt-bearing lawful slice with an explicitly non-inhabitable center.
2. A cross-channel transport criterion that preserves anchors, forbidden paths, authority ceilings, source-tense, receipts, and renarrow triggers rather than merely semantic similarity.
3. A terrain-shaped mixture-of-models formalism in which heterogeneous models propose morphisms but cannot confer permission, mutate constitutional state, or crown consensus.
4. A method for making high-value qualitative objects computational through path conditions, lifecycle, costs, receipts, and correction—without reducing their remainder to a scalar proxy.
5. A source-tense-aware engineering record and a falsifiable benchmark program, including exact symbolic reproduction of the announced Jacobian counterexample used as an external case study.
6. A reproducible paper-authoring receipt: each section was constrained by a machine-readable splat and linked into a hash chain before prose compilation.

Contents

1	Introduction: From Capability to Constitutional Geometry	6
1.1	The paper was authored through the mechanic it describes	7
1.2	Authorship and AI assistance	7
2	Claim Boundary and Evidence Classes	7
3	The Failure Class: When Local Validity Becomes Global Authority	8
3.1	Projection failures recur across domains	8
3.2	Collapse as topology loss	9
3.3	Why scalar correction is insufficient	9
4	Fractal Quivers of Quivers: The Parent Substrate	10
4.1	Classification is lawful traversal	10
4.2	Every facet is itself a field	11
4.3	Relationship to adjacent computational traditions	11
5	The Splat Primitive	11
5.1	The six facets	12
5.2	APO is forced as enforcement, conditional as expression	12
5.3	The working centroid	13
5.4	The minimum viable splat	13
6	Computing Around the Open Center	13
6.1	The center is not the centroid	14
6.2	A family of local centers over a founding center	14
6.3	Operational state without ontological closure	15
7	Cross-Channel Dehydration and Rehydration	15
7.1	Structural fidelity	16
7.2	Why natural language is a powerful exchange layer	16
7.3	Shape rehydration versus lexical retrieval	17
7.4	A cost model	17
8	Terrain-Shaped Mixtures of Models, Languages, and Places	18
8.1	Terrain is more than embedding space	18

8.2	Languages as model-local instruments	19
8.3	Place is part of the type	19
8.4	Separated convergence	19
8.5	A terrain-shaped model field in practice	19
9	Making Previously Noncomputational Objects Computational	20
9.1	Authority as geometry	21
9.2	Ambiguity as a held field	21
9.3	Dissonance as an absorber input	21
9.4	Trust as correction-path survival	21
9.5	Telos as a longitudinal hypothesis	22
9.6	Institutional memory as receipts over recall	22
9.7	The qualitative remainder	22
10	Longitudinal Learning by Receipts	22
10.1	Receipt-mediated visibility	23
10.2	Concentration as an anti-cheat signature	23
10.3	A maturity gate	24
10.4	Training multiple models without one model owning the terrain	24
11	Engineering Case Study: A Receipt-Bearing Federation	24
11.1	Clock, receipts, and the meaning of “700+ tics”	24
11.2	The conformance DAG refused a seamlessness claim	25
11.3	Shape-field rehydration	25
11.4	What the case study establishes	26
12	External Case Study: The Jacobian Lane	26
12.1	The exact finite certificate	26
12.2	How a perfect local sensor misses the failure	27
12.3	The Jacobian event rendered as a splat	27
12.4	What the event validates—and what it does not	28
13	Formal Properties and Design Lemmas	28
13.1	Anchor and exclusion preservation	29
13.2	Authority non-amplification	29
13.3	Receipt-carrying composition	29

13.4	Correction-path trust cannot be reduced to one slice	30
13.5	Collapse-path reduction	30
13.6	The guarantee/observation/hypothesis split	30
14	Evaluation Program: How the Paradigm Can Fail	31
14.1	Core research questions	31
14.2	Benchmark families	31
14.2.1	Cross-channel transport	31
14.2.2	Context laundering	31
14.2.3	Authority spoofing	32
14.2.4	Local-green/global-red systems	32
14.2.5	Longitudinal correction	32
14.2.6	Model-field separation	32
14.3	Metrics	32
14.4	Ablations	33
14.5	Falsifiers	33
15	Implementation Architecture: From Splat to Accountable Action	34
15.1	The stack	34
15.2	Reference algorithm	34
15.3	Schemas and receipts	36
15.4	Compare-and-swap and currentness	36
15.5	Absorbers and reversibility	37
15.6	Publication and interoperability	37
16	Limitations, Falsifiers, and the Research Agenda	37
16.1	Ontology cost and dialect risk	37
16.2	Boundary selection remains a human and institutional act	37
16.3	Complexity can create new bypasses	38
16.4	The center can be invoked performatively	38
16.5	The current evidence base is narrow	38
16.6	The Jacobian crosswalk may be overextended	38
16.7	Research agenda	38
17	Conclusion: A Constitutional Geometry for Compute	39

A	Notation	40
B	Reference Splat Schema	41
C	Exact Jacobian Verification	41
D	Federation Evidence Ledger	41
E	Paper Build Receipt	42
F	Submission and Review Checklist	42

1 Introduction: From Capability to Constitutional Geometry

The dominant story of AI progress is a story of model capability. More parameters, better data, longer contexts, faster inference, better tool use, stronger reasoning, and more capable agents increase what a system can do. That story is real. It is also incomplete. Once a model is connected to memory, tools, communication surfaces, filesystems, code execution, money, schedules, or delegated authority, the model no longer merely answers. It traverses. It reads, routes, edits, cites, invokes, stores, proposes, commits, escalates, and sometimes acts in another person's or institution's name.

At that point the governing question changes. The question is no longer only whether the model can produce a correct answer. It is whether the surrounding architecture makes the right motions possible, the wrong motions impossible or costly, the uncertain motions holdable, the irreversible motions separately authorized, and every material transition receiptable. Prompt-level rules can influence outputs, but they do not by themselves type authority, separate deliberation from disclosure, preserve source-tense, enforce standing, or make correction survive context loss. The failure is architectural before it is linguistic.

Prompted LLC publicly describes Ubiquity as a governance substrate for sovereign adaptive systems: software and logical infrastructure intended to increase AI-mediated capacity without collapsing agency, authorship, judgment, or meaningful contribution [Prompted LLC, 2026]. Its prior mathematical paper, *Fractal Quivers of Quivers*, frames models as bounded morphism-proposers inside a governed possibility topology, with classification defined by lawful traversal rather than a scalar label [Taylor, 2026]. This paper develops the compute primitive that sits between that broad topology and executable commitments: the *splat*.

A splat is a bounded, rehydratable representation of lawful problem shape. It retains enough positive structure to support action, enough negative structure to prevent premature collapse, enough provenance to reconstruct its ancestry, enough authority geometry to prevent fluent overreach, and enough unresolved state to be corrected later. It does not attempt to serialize the whole. Instead, it computes around a center that remains held open.

This design yields a particular inversion of the usual representation story. In standard machine learning, a representation is often judged by how compactly it captures information useful for prediction. In splat mechanics, a representation is also judged by what it refuses to erase: forbidden transitions, dissenting rays, unresolved contradictions, source-tense, authority ceilings, standing, and the distinction between a navigational centroid and the thing itself. Compression is valuable only if the boundary that keeps the object honest survives the move.

Load-bearing line

The frontier is not only better models. It is better constitutional geometry around models: a way to keep local competence, route plural intelligence, preserve contradiction, and refuse the crown.

The claim is intentionally strong and intentionally bounded. Splat mechanics is proposed as a new compute paradigm because it changes the primitive from *input mapped to output* to *governed possibility narrowed into lawful motion*. It does not follow that every splat implementation is safe, efficient, or superior. A bad ontology, hidden authority, missing receipts, stale runtime, or ornamental six-facet vocabulary can reproduce the same failure under more elaborate names. The architecture earns confidence only where its exclusions bind runtime and its receipts survive independent reconstruction.

1.1 The paper was authored through the mechanic it describes

This paper was not merely written *about* splats. Before prose compilation, each section was given a machine-readable working centroid, six-facet read, anchors, forbidden collapses, receipts, unresolved questions, and renarrow triggers. A validator checked all seventeen section splats and linked them into a SHA-256 chain. The source splat file has hash

c97f74ae123385cc5ecdeb62600c10789d0f9a6fc142a294e8283a341c62ebf5,

and the terminal section-ledger hash is

627ab53dd67128a6b2565e6d226bbb08850d53f28792ab820b0b80180bc97df7.

The ledger is a receipt of authoring geometry, not a certificate that the paper is true. Its function is narrower: it shows that the argument was constructed under explicit non-collapse constraints rather than retroactively decorated with them.

1.2 Authorship and AI assistance

Breyden E. Taylor is the author and bears responsibility for the thesis, terminology, claims, boundaries, and decision to publish. Homeskillet-S56 assisted with source synthesis, formal exposition, drafting, exact symbolic verification, artifact generation, typesetting, and build inspection. The Ubiquity Federation supplied prior artifacts and operational receipts. AI assistance is part of the production history. It is not mathematical or empirical evidence. Exact scripts, public sources, internal receipts, and independently reproducible observations carry the evidentiary load.

2 Claim Boundary and Evidence Classes

A governance-native paper should not force unlike evidence into one confident tense. This paper therefore distinguishes five evidence classes throughout.

Scope boundary

This paper does **not** claim: universal optimality; complete production hardening; a proof that a held-open center guarantees safety; a reduction of persons, institutions, ethics, or meaning to machine state; that the announced Jacobian counterexample proves Ubiquity; that one trainer pass establishes sustained learned terrain; or that internal engineering receipts substitute for external replication.

It does claim four narrower things.

First, the failure class is architectural. A system that permits a locally valid projection to silently acquire global authority is vulnerable even when the underlying model is intelligent and the local certificate is correct.

Second, the counter-shape can be expressed formally. A governed quiver with typed paths, anchors, apophatic ideals, absorbers, costs, receipts, standing, lifecycle, and center exclusion can make important collapse paths structurally inadmissible.

Table 1: Evidence classes used in this paper.

Class	Meaning	Permitted claim
Public	Published or publicly inspectable Prompted LLC surface.	“Publicly represented as” or “publicly documented.”
Internal receipt	Author-controlled engineering artifact with path, date, and hash in the bundle.	“Observed in the frozen engineering estate,” not independently replicated.
Author report	Current state reported by the author but not fully frozen into this release.	“The author reports,” never silently upgraded to public proof.
Exact reproduction	A finite calculation or artifact build reproduced by a supplied script.	“Verified exactly in this bundle” within the script’s claim boundary.
External announced	New public result with inspectable evidence but an incomplete ordinary review cycle.	“Announced,” “publicly checked,” or “as of the paper date.”

Third, a bounded lawful slice of that structure can be transported. When the carrier preserves anchors, exclusions, receipts, authority ceilings, and reconstruction conditions, problem shape can move across channels with higher structural fidelity than lexical or surface similarity alone.

Fourth, these claims are testable. The paper specifies falsifiers, conformance tests, blind cross-channel evaluations, adversarial context-selection tests, cost measures, and replication gates. If implementations cannot preserve the promised invariants under those tests, the paradigm has failed at the level that matters.

3 The Failure Class: When Local Validity Becomes Global Authority

Many computational systems are built around a useful projection: a score, label, embedding, proof obligation, benchmark, confidence value, policy, routing weight, or model response. Projections are not the problem. Computation requires them. Collapse begins when the projection’s jurisdiction disappears.

Let X denote a complex object or field and let

$$\pi_i : X \longrightarrow Y_i$$

be a projection useful for some task. A collapse-prone architecture behaves as though there exists an i for which

$$\pi_i(x) \text{ is locally useful} \implies \pi_i(x) \text{ is globally authoritative about } x.$$

The implication is usually unstated. It arrives operationally: the score controls the gate, the embedding decides identity, the model output changes the record, the locally green component is reported as globally complete, or a context is selected because it legalizes a desired action.

3.1 Projection failures recur across domains

The same topology appears under different names:

- A model’s confidence is treated as evidence of permission.
- Connected agents repeat one frame and are counted as independent convergence.
- A high semantic-similarity score is treated as proof that a boundary survived translation.
- A prompt-level refusal is treated as a runtime prohibition even though another path bypasses the prompt.
- A component test suite passes and the cross-repository system is called complete.
- A stated purpose is treated as authority over the behavior it is meant to constrain.
- A locally invertible map is assumed to be globally invertible.

None of these local artifacts must be false. The confidence can be calibrated, the agents can be capable, the embedding can be useful, the prompt can be well written, the component can be green, the purpose can be sincere, and the Jacobian can be exactly nonzero. The failure lies in promotion beyond jurisdiction.

3.2 Collapse as topology loss

A healthy decision field retains multiple rays under anchors. It supports refusal, dissent, review, correction, rollback, and renarrowing. A captured field progressively removes those alternatives until one frame becomes indistinguishable from truth, authority, care, urgency, or obligation. We call this *collapse* because the loss is structural: fewer lawful paths remain visible, but not because evidence eliminated them. They disappear because the representation or authority surface could not hold them.

For a path p , a working centroid z , and a held-open center $\odot$, define a coarse collapse predicate

$$\text{Collapse}(p, z) = [p \in \Omega] \vee [\text{Crown}(z, \odot)] \vee [\Gamma(p) = \emptyset] \vee [\text{ContextLaunder}(p)] \vee [\text{AnchorLoss}(p)]. \quad (1)$$

This is not yet a calibrated probability model. It is a typed failure surface. The empirical research program later asks how often these events occur under baseline and splat-governed systems.

Claim Boundary 3.1. *Equation (1) is a design-level indicator, not a claim that every form of systemic collapse has been enumerated or that each term can be measured without domain judgment.*

3.3 Why scalar correction is insufficient

One response is to add more scores: risk, uncertainty, fairness, confidence, cost, safety, and provenance. This helps, but it does not solve the authority problem. A vector of scores is still a projection, and aggregation silently introduces a sovereign function. If

$$R(x) = \sum_k w_k r_k(x),$$

then the weights w_k , the omitted dimensions, the update policy, and the threshold become constitutional decisions whether or not the interface calls them that.

Splat mechanics does not reject scores. It subordinates them. Scores may shape traversal pressure W_c ; they do not become truth. They can propose a route; they cannot grant standing, erase a forbidden path, satisfy a missing receipt, or inhabit the center.

4 Fractal Quivers of Quivers: The Parent Substrate

Splat mechanics is not the root architecture. It is one bounded computational lane inside a larger governed possibility topology. The parent object is the fractal, longitudinal, governed quiver-of-quivers developed in [Taylor, 2026]. The distinction matters because a splat can only be lawful relative to the anchors, authority, standing, lifecycle, costs, receipts, and telos of the field from which it is cut.

Definition 4.1 (Governed quiver). *For context c , a governed quiver is*

$$GQ_c = (V_c, E_c, s_c, t_c, \lambda_c, W_c, A_c, \Omega_c, B_c, \kappa_c, \Gamma_c, Z_c), \quad (2)$$

where V_c is the set of inhabitable vertices; E_c is the set of directed, typed edges; s_c, t_c are source and target maps; λ_c labels edge types; W_c carries traversal pressure but not truth; A_c contains anchors and invariants; Ω_c is an apophatic ideal of forbidden paths; B_c contains absorbers and bounded rewrites; κ_c prices coordination and irreversible movement; Γ_c is the receipt surface; and Z_c contains non-inhabitable center markers.

Typical vertices include entities, artifacts, claims, offices, models, duties, tools, states, and communication surfaces. Typical edges include **act**, **cite**, **inherit**, **route**, **expose**, **mutate**, **delegate**, **review**, **sync**, **train**, **freeze**, **fork**, **promote**, and **demote**. A graph records reachability. A governed quiver records whether that reachability is typed, lawful, current, authorized, reversible, priced, and receiptable.

Definition 4.2 (Governed quiver-of-quivers). *An estate-level object is*

$$\mathcal{Q}_c = (I_c, \{GQ_{i,c}\}_{i \in I_c}, \mathcal{M}_c, \Phi_c, \odot_c), \quad (3)$$

where I_c indexes subsystems or offices, $\mathcal{M}_c(i, j)$ contains admissible inter-quiver morphisms, Φ_c contains global compatibility constraints, and $\odot_c \in Z_c$ is the local held-open center.

A morphism

$$m : GQ_{i,c} \rightarrow GQ_{j,c}$$

may transport a claim, task, authority, receipt, model proposal, observation, training signal, or state transition. It is not lawful merely because it can be serialized. At minimum it must preserve the anchors and forbidden paths relevant to the move, satisfy authority and standing, and produce a receipt:

$$m(A_i) \subseteq A_j, \quad m(\Omega_i) \subseteq \Omega_j, \quad \Gamma_j(m(p)) \neq \emptyset. \quad (4)$$

Additional conditions may bind source-tense, lifecycle, currentness, cost, privacy, and parent telos.

4.1 Classification is lawful traversal

The longitudinal object is not a score. For observations or events e_i , let $R(e_i)$ be the governed record and let ω_i attenuate overrepresented strata without erasing them. Then

$$T_x(t) = \text{wcolim}_{i \leq t} (R(e_i), \omega_i) \in \text{GQuiv}. \quad (5)$$

The lawful traversal set is

$$\mathcal{L}_{x,c,t} = \{p \in \text{Path}(T_x(t) \hookrightarrow Q_c) : p \notin \Omega_{x,c,t}, \text{Auth}_c(p), \text{Standing}_c(p), \quad (6)$$

$$\text{Current}_c(p), \text{Lifecycle}_c(p), \Gamma_c(p) \neq \emptyset, t(p) \neq \odot_c\}. \quad (7)$$

The classifier is a rendering of permissions, obligations, refusals, holds, reviews, delegations, training actions, freezes, promotions, demotions, retirements, or receipts:

$$\mathcal{C}(x, c, t) = \Pi_c(\mathcal{L}_{x,c,t}). \quad (8)$$

The primitive is therefore not only $f : X \rightarrow Y$. It is a context- and history-sensitive computation of which motions remain lawful.

4.2 Every facet is itself a field

The six reading facets are

$$\Lambda = \{\text{KAT}, \text{APO}, \text{PAR}, \text{PLE}, \text{ENA}, \text{TEL}\}. \quad (9)$$

They are not six labels attached after the fact. Each is itself a quiver-of-quivers, and each can be refined at a finite facet address $\sigma \in \Lambda^*$. Thus KAT.APO.TEL and PAR.PLE.ENA identify different local reads, constraints, and routes. The repetition is fractal: the same grammar returns at each scale, but its local centroid, weights, exclusions, receipts, and telos can change.

This is why splat mechanics is not simple recursion. Recursion repeats an operation. Fractal traversal reuses a grammar while allowing the local geometry to change by address.

4.3 Relationship to adjacent computational traditions

The substrate touches several established traditions without collapsing into any one of them. Abstract argumentation represents attacks and acceptability among arguments [Dung, 1995]. Provenance standards represent derivation and responsibility across systems [World Wide Web Consortium, 2013]. Applied category theory studies compositional structure and interfaces [Fong and Spivak, 2019, Lane, 1998]. Conceptual spaces and hyperdimensional computing provide geometric and distributed representations [Gärdenfors, 2000, Kanerva, 2009]. Safety shielding constrains actions independently of a learner’s preferred policy [Alshiekh et al., 2018]. Model cards and datasheets improve disclosure about models and datasets [Mitchell et al., 2019, Gebru et al., 2021].

Splat mechanics binds concerns that these traditions usually hold separately: compositional path structure, negative constraints, authority, standing, lifecycle, cost, provenance, purpose, open-center exclusion, and longitudinal correction. The claim is not that no prior field contains related pieces. The claim is that their particular cross-binding creates a distinct computational primitive.

5 The Splat Primitive

A splat is the current, bounded, conformation-specific narrowing through which a governed field becomes portable and actionable. It is a slice, not the whole; a carrier, not the center; a current reference, not final doctrine.

Definition 5.1 (Governed splat). *For object or field x in context c at time t , a governed splat is*

$$\mathbf{S}_{x,c,t} = (Q^S, A^S, \Omega^S, B^S, \kappa^S, \Gamma^S, \Lambda^S, z^S, \odot_c, U^S, \mathcal{R}^S), \quad (10)$$

where Q^S is the admitted local subquiver; A^S its active anchors; Ω^S its forbidden-path ideal; B^S its available absorbers; κ^S its cost surface; Γ^S its receipts and provenance; Λ^S its six-facet cross-binding; z^S a movable working centroid; $\odot_c$ the non-inhabitable local center marker; U^S the unresolved or suspended set; and $\mathcal{R}^S$ the renarrow triggers.

The representation is deliberately richer than a vector and deliberately smaller than the source estate. It should be compact enough to carry and complete enough to refuse the most dangerous distortions.

5.1 The six facets

For an event, artifact, question, claim, behavior, or proposed action e , each facet contributes a local quiver $L_\ell(e)$:

$$L_\ell(e) \in \mathbf{GQuiv}, \quad \ell \in \Lambda. \quad (11)$$

The record is cross-bound rather than concatenated:

$$R(e) = \operatorname{colim}_{\ell \in \Lambda} L_\ell(e) / \sim_{\text{cross}}. \quad (12)$$

The quotient $\sim_{\text{cross}}$ prevents the facets from remaining decorative notes. A claim that “this is a runtime hook” under KAT, “must not become doctrine” under APO, “mutation authority belongs to review” under PAR, “requires a current sync receipt” under PLE, “can look installed while stale” under ENA, and “serves runtime parity” under TEL becomes one directed structure. The meaning lies in how these constraints alter the available paths.

Table 2: The six facets as computational functions.

Facet	Natural question	Computational function
KAT	What is positively present?	Names the current mechanic, object, state, source-tense, and anchor truth that can be asserted now.
APO	What must not collapse?	Removes forbidden paths, false equivalences, premature crowns, illegal imports, and unsafe context shifts.
PAR	What is present but not centered?	Surfaces hidden authority, burden, dignity, visibility, standing, and unannounced power.
PLE	What complement is required?	Names interfaces, activation conditions, receipts, handoffs, runtime support, and missing halves.
ENA	How does it fail or invert?	Models misuse, overuse, underuse, staleness, self-defeat, and apparently-correct wrong pulls.
TEL	What does the motion serve?	Binds local subtelos to parent purpose, reversibility, and longitudinal behavioral direction.

5.2 APO is forced as enforcement, conditional as expression

Apophatic constraint is always active as a gate. Apophatic language is not always the best way to render every other facet. Let

$$\Omega_{x,c,t} = \text{Ideal} \left(\Omega_{x,c,t}^{\text{APO}} \cup \bigcup_{\ell \neq \text{APO}} \Omega_{x,c,t}^{\ell, \text{internal}} \cup \Omega_{\text{ctx}} \cup \Omega_{\text{center}} \right). \quad (13)$$

Then

$$p \in \Omega_{x,c,t} \implies p \notin \mathcal{L}_{x,c,t}. \quad (14)$$

The choice to render a non-APO facet through its own apophatic substructure is controlled separately:

$$\varepsilon_{\ell \leftarrow \text{APO}}(x, c, t) = 1 \iff \text{profit}_{\ell, \text{APO}} - \text{complexity}_{\ell, \text{APO}} > \tau_c. \quad (15)$$

Plainly: the gate stays even when the speech becomes positive. APO prevents the other five facets from lying too early; it does not require them to remain linguistically negative after anchored truth becomes available.

5.3 The working centroid

A working centroid summarizes the navigational core of the lawful local field:

$$z_{x,c,t}^\sigma = \text{centroid}_\mu \left(\text{Core}(\text{Path}(Q_{x,c,t}^\sigma) \setminus \Omega_{x,c,t}^\sigma) \right). \quad (16)$$

It is computed, facet-relative, gauge-relative, movable, and useful. It may be represented as a vector, graph core, prototype, basis, set of poles, or structured natural-language object. What matters is not the data type but its constitutional status: it is a coordinate into the field, not the whole field and not the center.

5.4 The minimum viable splat

A splat must contain enough information to answer eight questions:

1. What is currently being asserted, and in what source-tense?
2. What anchors must survive any translation or action?
3. What paths are forbidden regardless of model preference?
4. Who has authority and standing to propose, approve, mutate, or receive?
5. What remains unresolved, suspended, or contingent?
6. What receipts establish currentness and ancestry?
7. What failure inversions and absorbers are available?
8. What event would force the slice to be recomputed?

A summary can omit these and still be readable. A splat cannot omit them and remain lawful.

6 Computing Around the Open Center

The open center is the load-bearing difference between a splat and a totalizing representation. It is not a blank to be filled later, a confidence interval, an unknown class, or a poetic gesture. It is a typed exclusion from the inhabitable state space.

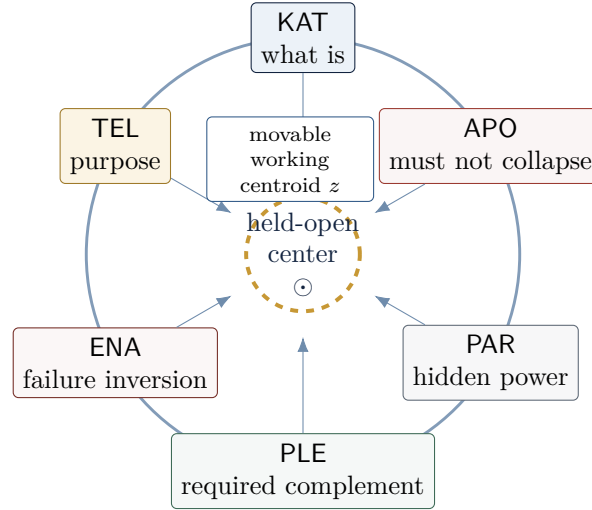
Definition 6.1 (Center sort separation). *Let V_c be the inhabitable vertices of context c and Z_c the center-marker sort. Require*

$$V_c \cap Z_c = \emptyset, \quad \odot_c \in Z_c. \quad (17)$$

Edges type only over inhabitable vertices:

$$s_c, t_c : E_c \rightarrow V_c. \quad (18)$$

Therefore no edge or path can begin or end at $\odot_c$.



The perimeter computes a lawful local field. The center remains outside the state space.

Figure 1: The splat perimeter. The six facets constrain and complete a working centroid while the held-open center is explicitly non-inhabitable.

The stronger runtime condition is

$$\odot, \odot_c \notin \Theta_{\text{trainable}} \cup \Theta_{\text{routeable}} \cup \Theta_{\text{serializable as mutable state}} \cup \Theta_{\text{objective}} \cup \Theta_{\text{model output}}. \quad (19)$$

The center is not merely protected from mutation. It is unavailable as a target. A model cannot predict it, an optimizer cannot minimize distance to it, a router cannot send work into it, and a database cannot save it as the current truth.

6.1 The center is not the centroid

The distinction is categorical:

	Working centroid $z_{x,c,t}^\sigma$	Held-open center $\odot_c$
Computability	computed	not computed
Role	navigation	anti-capture boundary
Dependence	context-, time-, and gauge-relative	context-indexed over a founding center
Movement	movable	never traversed
Training	may be estimated	not a trainable target
Serialization	may be receipted	not mutable state
Authority	provisional coordinate	confers no answer

The center does not answer. It prevents answers from becoming crowns.

6.2 A family of local centers over a founding center

Local center exclusion is insufficient if an actor can choose a friendlier context. Let $\odot_c$ be the local center marker and $\odot$ the founding center. Require a projection

$$\pi_c(\odot_c) = \odot. \quad (20)$$

For a context transition $f : c \rightarrow d$ with induced transport $f_* : GQ_c \rightarrow GQ_d$, admissibility requires preservation of the founding center and founding telos:

$$\pi_d(f_*(\odot_c)) = \odot, \quad \text{Telos}(f_*(\odot_c)) = \text{Telos}(\odot). \quad (21)$$

A context may change the local gauge. It may not select a new founding reality merely to legalize a path forbidden in the old one. We call that failure *context laundering*.

Design Lemma 6.2 (Center noncapture by construction). *If center markers are disjoint from inhabitable vertices as in (17), all edges type only over inhabitable vertices, and the runtime excludes center markers as in (19), then no well-typed traversal, model output, training objective, or mutable state can inhabit the center.*

Proof. Every well-typed path is composed of edges whose sources and targets lie in V_c . Since $V_c \cap Z_c = \emptyset$, no path reaches $\odot_c$. The runtime exclusion separately removes the center from trainable, routeable, objective, output, and mutable-state domains. An attempted center mention is therefore either descriptive metadata or an ill-typed center-capture event requiring refusal and receipt; it cannot be ordinary state transition. $\square$

This lemma is modest but important. It does not prove that the surrounding system is safe. It proves that one class of capture—the conversion of a representation into an inhabitable crown—is structurally unavailable when the premises actually bind runtime.

6.3 Operational state without ontological closure

The apparent tradeoff is that refusing to compute the center might make the system indecisive. Splat mechanics takes the opposite position. Operational state comes from a dense, bounded working centroid and a rich perimeter, not from pretending the entire object has been captured. A system can know enough to route, refuse, stage, ask, execute, verify, or rollback while explicitly retaining what it does not know and what it is not permitted to become.

This is especially valuable in domains where the object includes persons, institutions, evolving projects, ethical duties, or open scientific questions. The system can represent current commitments, constraints, evidence, and action paths without reducing the person, institution, or field to the current record.

The open-center proposition

The open center does not make computation less concrete. It relocates concreteness from a totalizing state estimate to a governed perimeter of lawful motion. The system acts on what is receipted, holds what is unresolved, and keeps the remainder unavailable for coronation.

7 Cross-Channel Dehydration and Rehydration

The economic promise of splat mechanics is not that every rich source object can be copied everywhere. It is that the structural shape required for lawful continuation can be dehydrated into a compact carrier and rehydrated in a target channel without transporting the entire source surface.

Let X_i and Y_j be source and target representation spaces: natural language, code, graphs, telemetry, images, database records, model activations, institutional procedures, or human-readable policy. A splat

transport is

$$X_i \xrightarrow{E_i} S_{x,c,t} \xrightarrow{D_j} Y_j, \quad (22)$$

where E_i is a source encoder and D_j is a target decoder. The carrier S is not required to preserve every token, coordinate, or aesthetic property. It must preserve the structure that determines lawful motion.

7.1 Structural fidelity

Semantic similarity is one component of fidelity, not its sovereign measure. Define a transport morphism $m : S_i \rightarrow S_j$. A minimum admissibility condition is

$$m(A_i) \subseteq A_j, \quad (23)$$

$$m(\Omega_i) \subseteq \Omega_j, \quad (24)$$

$$\text{AuthCeiling}_j(m(p)) \leq \text{AuthCeiling}_i(p), \quad (25)$$

$$\Gamma_j(m(p)) \neq \emptyset, \quad (26)$$

$$\text{SourceTense}_j(m(x)) = \text{TransportedTense}_i(x). \quad (27)$$

The target may add stricter constraints. It may not silently delete source anchors, legalize forbidden paths, increase authority, or upgrade a reported claim into runtime fact.

A practical fidelity vector is

$$\mathbf{F}(m) = (F_A, F_\Omega, F_R, F_\Gamma, F_T, F_U, F_C), \quad (28)$$

where:

- F_A measures anchor retention;
- F_Ω measures forbidden-path retention;
- F_R measures relation and dependency preservation;
- F_Γ measures receipt and provenance completeness;
- F_T measures source-tense and lifecycle preservation;
- F_U measures preservation of unresolved and suspended state;
- F_C measures counterfactual behavior under known collapse-zone probes.

A lexical or embedding similarity score may be appended, but it cannot compensate for a zero in an anchor or exclusion dimension.

Definition 7.1 (Lawful cross-channel transport). *A transport m is lawful when every required hard dimension of $\mathbf{F}(m)$ exceeds its domain threshold, the target authority does not increase without a separate grant, and the target emits a receipt sufficient to reconstruct the source splat and transformation version.*

7.2 Why natural language is a powerful exchange layer

Natural language can carry mutually valid frames, counterfactuals, exceptions, source-tense, purpose, uncertainty, and relations that are expensive to encode in one rigid formalism. The Ubiquity methods inventory describes cross-domain transport in which ecological, immunological, thermodynamic, and semiotic mechanics are carried through language and retyped onto governance surfaces [Ubiquity Federation, 2026b]. This is not merely metaphor when the imported mechanic changes routing, thresholds, retrieval, or refusal behavior and is tested against the destination’s anchors.

Language has two advantages as an exchange layer. First, it is broadly decodable by heterogeneous models and humans. Second, it can preserve conditionality without forcing early branch selection. Its disadvantage is equally important: fluent restatement can create the appearance of fidelity while changing authority or deleting a negative boundary. A language carrier therefore requires the same morphism tests as code or schema transport.

7.3 Shape rehydration versus lexical retrieval

An internal Ubiquity candidate proof tested a shape-field rehydration mechanism in which a splat-Mahalanobis nearest-neighbor read was retyped onto a six-ray terrain. A lexical grep misrouted a case because it matched the word used to describe rehydration; the shape-based method preserved the governing doctrine across a full corpus. Two independent blind encoders, neither shown the proof construction, produced encodings under which the same nearest-neighbor decision survived across four test families. The result was correctly retained as a candidate proof rather than promoted directly to doctrine [Ubiquity Federation, 2026c].

The significance is not the specific distance function. It is the separation between *surface token* and *governing shape*. A transport mechanism should be robust to paraphrase and encoding style while remaining sensitive to changes in anchors, collapse zones, authority, and telos.

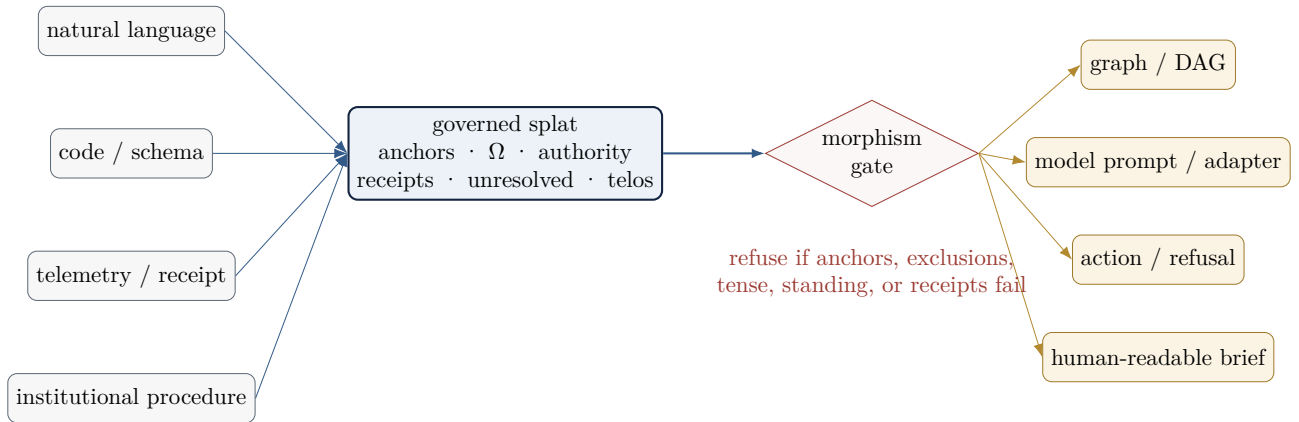


Figure 2: Cross-channel splat transport. The carrier preserves lawful problem shape; a target decoder does not gain authority merely by being able to render it.

7.4 A cost model

Let $C_{\text{full}}(X_i, Y_j)$ be the cost of transporting and reconstructing the full source surface, and let

$$C_S = C_E + C_{\text{carrier}} + C_{\text{gate}} + C_D + C_{\Gamma} \quad (29)$$

be the cost of encoding, carrying, gating, decoding, and receipting a splat. The method is economically valuable when

$$C_S < C_{\text{full}} \quad \text{and} \quad \mathbf{F}(m) \succeq \tau \quad (30)$$

for hard fidelity thresholds τ . This is not guaranteed. A high-stakes domain may require a splat so rich that full-state transport is cheaper or safer. The point is to make the tradeoff explicit rather than letting cheap surface similarity masquerade as adequate reconstruction.

Design Lemma 7.2 (Compositional preservation). *Let $m_1 : S_1 \rightarrow S_2$ and $m_2 : S_2 \rightarrow S_3$ preserve the relevant anchors and forbidden paths, never increase authority without a separate receipt, and emit reconstructible receipts. Then $m_2 \circ m_1$ preserves those properties provided the target domains interpret the preserved symbols compatibly.*

Proof. Anchor and forbidden-path preservation follow by set inclusion under composition. Authority nonincrease is transitive when each stage’s ceiling is monotone. Receipt reconstruction follows by chaining the predecessor hashes and transformation identifiers. The compatibility premise is load-bearing: identical labels with different destination semantics can break preservation despite syntactic inclusion. $\square$

8 Terrain-Shaped Mixtures of Models, Languages, and Places

Modern systems increasingly combine models. Mixture-of-experts architectures sparsely activate parameters inside a model [Fedus et al., 2022]; routers select among stronger and weaker models to balance quality and cost [Ong et al., 2024, Chen et al., 2023]; mixture-of-agents systems layer multiple LLM outputs to improve generation [Wang et al., 2024]. These approaches demonstrate that heterogeneous model fields can outperform one undifferentiated model or reduce cost. They do not, by themselves, solve constitutional authority.

Splat mechanics treats model heterogeneity as one ray inside a larger terrain. Different models can live in different locations, languages, access domains, cost regimes, privacy zones, offices, and lifecycle states. Their outputs are not votes that acquire truth by accumulation. They are typed proposals whose admissibility depends on the current splat.

Definition 8.1 (Bounded morphism-proposer). *A model or model-bound office is*

$$M_j = (X_j, Y_j, \phi_j, \psi_j, \alpha_j, \rho_j), \quad (31)$$

where X_j is its admissible input domain; Y_j its proposal codomain; ϕ_j its read map into a governed field; ψ_j its proposal map into candidate paths; α_j an authority predicate; and ρ_j its receipt requirement.

A model is active only when it is current, authorized for the input and context, and does not propose exclusively forbidden paths:

$$j \in J(x, c, t) \iff \alpha_j(x, c, t) = 1 \wedge x \in X_j \wedge \text{Current}_c(M_j) \wedge \psi_j(\phi_j(x)) \not\subseteq \Omega_c. \quad (32)$$

The model proposal field is

$$\mathcal{M}_{x,c,t} = \text{colim}_{j \in J(x,c,t)} \psi_j(\phi_j(x)). \quad (33)$$

This is a constrained colimit, not a softmax and not a majority vote. It combines proposals that survive type and authority checks; governance later admits, refuses, holds, stages, delegates, or asks for verification.

8.1 Terrain is more than embedding space

A terrain is a field of positions, pressures, collapse zones, costs, standing, authority, memory, and possible motion. An embedding may contribute the geometric coordinate. The terrain adds constitutional structure. For model j and path p , a routing energy can be written schematically as

$$E_j(p \mid S) = d_{\text{shape}}(\phi_j(x), z^S) + \lambda_\kappa \kappa_j(p) + \lambda_\Delta \Delta_{\text{current}}(M_j) \quad (34)$$

$$+ \lambda_\Omega \mathbf{1}[p \in \Omega] + \lambda_A \text{AnchorLoss}(p) + \lambda_R \text{ReceiptDeficit}(p) + \lambda_S \text{StandingViolation}(p). \quad (35)$$

A router may use E_j to rank candidates, but hard violations remain gates, not merely large penalties. Otherwise sufficient model confidence or economic value can buy passage through a supposedly forbidden boundary.

8.2 Languages as model-local instruments

A mixture of language models is also a mixture of languages in a broader sense: programming languages, natural languages, mathematical notations, institutional vocabularies, schemas, and visual grammars. Each language makes some relations cheap to express and others difficult. Natural language holds conditional tension; a type system enforces interfaces; a proof assistant checks formal derivations; a DAG exposes dependency; a database preserves current state; a visualization reveals shape; a human narrative can preserve lived stakes that a schema would flatten.

Splat mechanics does not demand one universal encoding. It uses a shared lawful grammar so each lane can express the same problem through its strongest instrument while preserving cross-lane anchors and exclusions. A model operating in code may propose a patch; a model operating in natural language may surface an authority ambiguity; a graph ranker may identify a neighboring doctrine; a local embedding model may find a terrain match; a remote frontier model may generate a proof strategy; a human may supply the telic and institutional judgment that none of them owns.

8.3 Place is part of the type

Compute location matters. A local model, a cloud model, a sandboxed tool, an office with private memory, and a public model endpoint do not inhabit the same authority or disclosure surface. Place therefore enters the context rather than remaining infrastructure trivia:

$$c = (\text{domain, time, place, standing, authority, lifecycle, telos}). \quad (36)$$

The same proposal can be lawful from one place and unlawful from another because the actor's standing, data access, confidentiality obligations, or receipt capabilities differ. This is not arbitrary relativism. Local contexts orbit a founding center and common anchors; context selection cannot legalize a path forbidden at the root.

8.4 Separated convergence

Agreement among connected agents can arise from shared context, shared contamination, or deference. Let $a_1, \dots, a_n$ be readers with approximately separated contexts and distinct mandates. If $I(a_i; a_j) \approx 0$ for $i \neq j$, then

$$\text{Conv}_{\text{sep}}(x) = \bigcap_i R_i(x) \quad (37)$$

is stronger evidence than connected echo, although it still does not create authority. The point is not that independence can be perfectly achieved. It is that dependence is a first-class variable rather than hidden behind a count of agreeing outputs.

8.5 A terrain-shaped model field in practice

The Ubiquity training stack described in internal receipts is intentionally layered: deterministic floor, local embedding model, cloud comparison, graph ranker, triplet reranker, and LoRA last. The sequence

Table 3: Splat-governed model fields compared with adjacent mixture patterns.

Pattern	Primary optimization	Splat-governed distinction
Mixture of Experts	Sparse parameter activation and scale efficiency.	Expert activation remains subordinate to hard anchors, authority, standing, and receipts.
RouteLLM / cascades	Cost-quality routing among models.	Cost and quality shape pressure; they cannot legalize forbidden paths or grant mutation rights.
Mixture of Agents	Aggregate or layer multiple model outputs.	Connected agreement is weak evidence; separated convergence and mandate diversity are tracked explicitly.
Ensemble voting	Reduce variance through consensus.	Votes remain proposals. Minority refusal can be load-bearing when it names a collapse zone.
Tool orchestration	Select tools by task.	Tool capability and tool authority are separate typed relations.
Splat model field	Lawful motion under terrain and governance.	The output is a path disposition with receipts, not merely a synthesized answer.

places behavior-conditioned terrain and retrieval before parameter adaptation. The author additionally reports multiple model artifacts and model classes in active training and evaluation across local and remote compute. This paper treats that current claim as author report, not as a public benchmark.

The principle is broader than the specific stack: the model should be trained by movement through terrain under receipts, not trained to believe the lane’s declared identity. A local model may learn proximity, a graph model may learn structure, and a language model may learn expression, while constitutional anchors remain outside the reach of any one training objective.

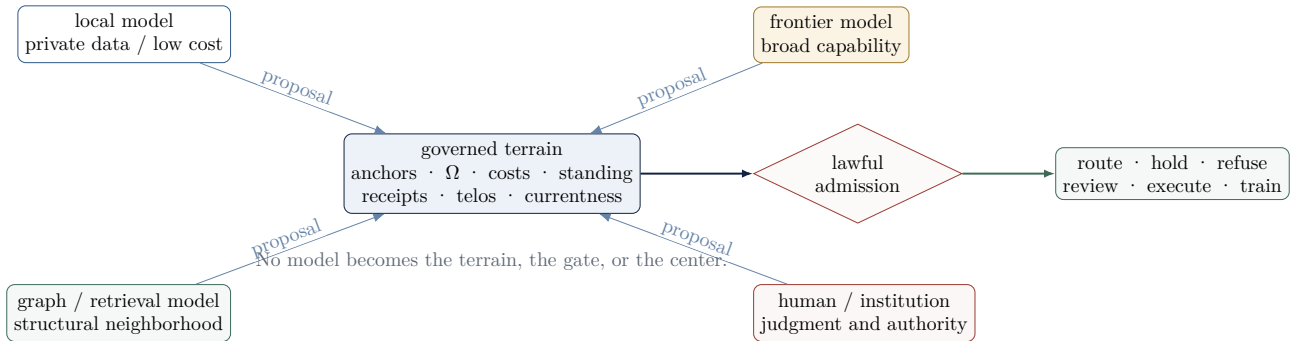


Figure 3: A terrain-shaped mixture of models, languages, and places. Heterogeneous actors propose; the governed field admits lawful motion.

9 Making Previously Noncomputational Objects Computational

Some of the most important constraints in real systems are often described as “soft”: authority, purpose, ambiguity, dignity, burden, trust, institutional memory, reversibility, and dissonance. They are soft only

when the architecture lacks primitives for them. Splat mechanics makes them computational by giving them observable path consequences, not by pretending to reduce their full meaning to a number.

9.1 Authority as geometry

Authority is not tone or model capability. It is a relation among actor, action, object, context, and time:

$$\text{May}(a, \alpha, o, c, t) \in \{\text{propose, review, approve, execute, mutate, receive}\}. \quad (38)$$

A proposal from a highly capable model can be structurally well formed and still lack mutation authority. Conversely, a low-capability component may hold narrow write authority because it is deterministic, auditable, and scoped. The architecture becomes computational when these relations are typed and enforced at the edge where motion occurs.

9.2 Ambiguity as a held field

Conventional pipelines often treat ambiguity as missing information to be resolved. Splat mechanics can retain multiple conditional rays:

$$\mathcal{H}_t = \{h_1, \dots, h_k\} \quad (39)$$

with each hypothesis linked to supporting evidence, exclusions, expected observations, authority implications, and renarrow triggers. An apophatic observation updates by elimination:

$$\mathcal{H}_{t+1} = \mathcal{H}_t \setminus \{h \in \mathcal{H}_t : h \text{ violates observed boundary behavior}\}. \quad (40)$$

The system computes with ambiguity by routing differently under different surviving hypotheses, not by forcing a premature positive identity.

9.3 Dissonance as an absorber input

Dissonance is not automatically error. It is a signal that two or more constraints cannot be jointly satisfied under the current slice. An absorber is a partial path rewrite

$$b \in B_c, \quad b : \text{Path}(Q_c) \rightharpoonup \text{Path}(Q_c) \times \Gamma_c. \quad (41)$$

It can map an unsafe path to reviewed path plus receipt, an overbroad path to scoped path plus receipt, a stale path to hold plus drift receipt, or a runaway path to freeze plus containment receipt. The system remains resilient when contradiction produces structured correction rather than denial or crash.

9.4 Trust as correction-path survival

Trust is not first-pass correctness, model confidence, or a single successful outcome. Let p be a longitudinal path. Define its trust object as

$$\text{TrustQuiver}(p, c) = (Q_p, \Gamma_p, B_p, \Omega_p, A_p, \odot), \quad (42)$$

where the object records the path taken, receipts, available absorbers, forbidden alternatives, preserved anchors, and the center it did not target. A scalar projection

$$d_{\text{trust}} : \text{TrustQuiver} \rightarrow [0, 1]$$

may support dashboards, but the object of trust is the correction-bearing path.

A system is trustworthy to the extent that it can show how a claim moved, what it could not become, where correction entered, who held authority, which falsifiers it survived, what receipt closed the loop, and what remains unresolved.

9.5 Telos as a longitudinal hypothesis

A stated purpose is evidence, not authority. Let T^* be latent true telos, T_{expr} expressed telos, T_{beh} behavioral direction, and T_{corr} correction direction. The estimate

$$\hat{T}_t = \text{RobustCore}(T_{\text{beh}, \leq t}, T_{\text{corr}, \leq t}, T_{\text{expr}, \leq t}, F_{\leq t}, \Gamma_{\leq t}) \quad (43)$$

weights expressed purpose more heavily only when it converges with behavior and correction under falsifiers. Trust binds only to identified telos:

$$\text{BindTrust}(u, T, t) = 1 \iff T \in \text{IdentifiedTeloi}_u(t). \quad (44)$$

This makes purpose computational without allowing fluent self-description to certify itself.

9.6 Institutional memory as receipts over recall

A language model can reconstruct a plausible history. A receipt persists. Provenance systems such as PROV-O make derivation interoperable [World Wide Web Consortium, 2013]; splat mechanics adds currentness, authority, forbidden paths, and correction triggers to the same concern. Memory becomes governance-bearing when a future actor can distinguish what was seen, told, staged, executed, verified, rolled back, superseded, or terminalized.

9.7 The qualitative remainder

Operationalization is not exhaustion. A splat can make a dignity burden visible as a constraint, standing relation, refusal, review requirement, or irreversible cost. It does not therefore contain dignity. It can represent the institutional consequences of a purpose without containing the full purpose. This remainder is protected by the open center.

Operationalization without reduction

A construct becomes computational when it changes typed motion, evidence requirements, cost, or correction. It need not become a scalar, and the architecture must preserve the difference between the operational proxy and the construct it serves.

10 Longitudinal Learning by Receipts

A terrain-shaped system can learn without immediately changing base-model parameters. Retrieval weights, graph edges, methylation, activation thresholds, route preferences, and local centroids can move as the system accumulates receipted behavior. This resembles Hebbian association in the broad sense—what repeatedly fires together becomes easier to retrieve together—but the update occurs in governed selection and memory rather than unbounded backpropagation.

The central integrity problem is self-flattery. A system can declare that it is computationally mature, add the desired words to its lane description, and then observe that its embeddings moved toward the declared identity. That is not learning. It is target leakage.

10.1 Receipt-mediated visibility

Let B_t be the behavioral receipt corpus and D_t the declared-lane or aspirational surface. Require source disjointness:

$$B_t \cap D_t = \emptyset \quad \text{at the evidence-construction layer.} \quad (45)$$

A learned office position is computed from behavior made visible through receipts:

$$z_t^{\text{learned}} = \text{LearnPosition}(B_{\leq t}), \quad (46)$$

while the declared lane produces a frontier or candidate direction

$$\mathcal{F}_t^{\text{declared}} = \text{Frontier}(D_t). \quad (47)$$

The learner may move toward the frontier only if behavior, not declaration, changes the learned position.

Code and commits are not automatically behavior receipts. They may prove that an artifact exists; they do not necessarily show that an office used it, survived correction, or earned a new standing. The trainer therefore requires each material slice to emit an auditable trace of what it did.

10.2 Concentration as an anti-cheat signature

Uniform movement across the whole frontier can arise from generic activity, lexical stuffing, or lane-weight leakage. Genuine behavior should move the patterns the behavior actually touched more than untouched controls.

An internal F2.2 trainer pass at tic 343 measured affinity before and after a concrete runtime build. The result is reported in Table 4 and preserved in the frozen artifact [Ubiquity Federation, 2026d].

Table 4: Pattern-specific affinity movement in the first F2.2 trainer pass. The result is an internal engineering observation with $n = 1$, not sustained-terrain proof.

Pattern	Affinity delta
schema-reject	+0.0175
sovereign-lane-TTL	+0.0171
dispatch-preflight	+0.0142
condition-stable-task_id	+0.0135
untouched tournament-lattice control	+0.0013

The four touched patterns moved roughly three to four times more than the untouched control. The source-disjoint false-flattery guard remained passing. The observation supports a specific claim: the learner detected behavior-specific motion rather than only whole-frontier warming.

Engineering Observation 10.1 (Correct withholding). *The system did not promote `compute` → `earned` after the single pass. The map recorded the direction but retained the claim at tentative/spec rung. In this setting, withholding is evidence of integrity: a learner that rewrites its earned center after one favorable event is easily flattered.*

10.3 A maturity gate

Let P_t be the posterior that a candidate terrain relation is earned and let n_{conf} count independent conformations under relevant falsifiers. Promotion requires

$$P_t > P_{\min}, \quad n_{\text{conf}} \geq n_{\min}, \quad \text{ControlLeak} = 0, \quad \text{ReceiptCoverage} \geq \tau_{\Gamma}. \quad (48)$$

The exact thresholds are domain-specific. The invariant is that one impressive slice cannot silently become longitudinal truth.

10.4 Training multiple models without one model owning the terrain

A terrain-learning stack can distribute functions across models:

- deterministic rules establish a floor and obvious refusals;
- local embeddings estimate shape proximity under private data constraints;
- graph rankers identify structural neighborhoods and dependencies;
- triplet or contrastive learners sharpen distinctions among nearby lawful and collapse-prone states;
- language models translate, propose, explain, and generate candidate morphisms;
- LoRA or other parameter adaptation occurs only after the training signal has earned maturity.

No layer owns the constitutional truth. The terrain is reconstructed from receipts and tested by behavior. Parameter learning is one possible expression of the field, not its sovereign source.

11 Engineering Case Study: A Receipt-Bearing Federation

Ubiquity has been developed inside a live software federation rather than as a paper-only design. The estate includes offices, councils, tics, receipts, model adapters, memory and retrieval systems, governance ledgers, training surfaces, conformance DAGs, and execution-oriented artifacts. The public State of the Federation at tic 516 names nineteen offices and explicitly describes a center the apparatus is forbidden to strike [Ubiquity Federation and ent_homeskillet, 2026]. The public site also presents the FQoQ formal paper and its companions as load-bearing surfaces [Prompted LLC, 2026].

This is useful evidence because the architecture has encountered real friction: stale state, incorrect promotion, cross-repository drift, false local green, missing emitters, authority confusion, source-tense collapse, lexical retrieval errors, and concurrency hazards. It is not a controlled benchmark. The originating estate can be coherent and still be idiosyncratic, self-confirming, or difficult to reproduce elsewhere.

11.1 Clock, receipts, and the meaning of “700+ tics”

A *tic* is the canonical governance clock, distinct from subordinate work units. Time is authored before work so temporal provenance does not depend on throughput. The author reports that the live federation had surpassed 700 tics by July 2026. The frozen evidence used for this paper does not collapse all nearby counts into that claim:

- the publicly inspectable federation address is tic 516;
- internal operating transcripts in the frozen library reach tic 622;
- a routing-decision ledger visible in those transcripts contained 701 rows by tic 621;
- 701 decision rows are not 701 tics, and the later live tic count is author report unless separately released.

This distinction is not clerical. It is the method: row count, tic count, public address, current runtime, and author report are different source-tenses and must remain different.

11.2 The conformance DAG refused a seamlessness claim

A July 2026 conformance audit examined a temporal-splat runtime across multiple repositories. Several component-local tests and continuous-integration runs were green, but the audit concluded that the bundle was not yet one current, proven federation. It preserved a lifecycle chain:

$$\begin{aligned}
 &\text{review-green} \neq \text{landed} \neq \text{registered-current} \\
 &\quad \neq \text{selector-eligible} \neq \text{activated} \\
 &\quad \neq \text{canonically absorbed} \neq \text{terminalized}.
 \end{aligned} \tag{49}$$

The artifact identified open joins among proposal validation, runtime input, execution envelopes, authoritative emitters, canary acquisition, selectors, cross-repository maps, and canonical absorption [Ubiquity Federation, 2026a]. The important result is not that the system had gaps. Every developing system has gaps. The important result is that local green was not allowed to impersonate global completion.

Design Lemma 11.1 (No promotion by local green). *If lifecycle states are typed as separate dimensions and readiness requires their conjunction, then success in any proper subset of dimensions cannot derive terminal readiness.*

Proof. Let readiness be

$$R = C \wedge P \wedge K \wedge D \wedge A \wedge G, \tag{50}$$

where C is admitted covenant, P complete projection, K current conformation, D satisfied dependencies, A available runtime, and G absence of active gates. A local test establishes at most a subset. Without all conjuncts, R is false. The result is trivial logically and often violated operationally when the dimensions are flattened into one status string. $\square$

11.3 Shape-field rehydration

The shape-field experiment described earlier provides a second engineering signal. A lexical retrieval method matched a shared word and selected the wrong governing doctrine. A shape-based nearest-neighbor read, with apophatic vetoes, selected the intended doctrine. The result survived two independent blind encodings and remained stable across four slice families. Because the mechanism had not completed every governance gate, it remained a candidate rather than a canonical primitive.

This is a small example of the broader claim: a splat can move through different encoders with higher problem-pattern fidelity than a token search, but only if the target preserves the relevant dimensions and the result survives blind reconstruction.

11.4 What the case study establishes

The engineering estate supports five bounded observations:

1. The vocabulary has executable counterparts: paths, gates, receipts, costs, absorbers, source-tense, currentness, and lifecycle are represented in artifacts and tests rather than only prose.
2. The system can refuse its own completion narrative when cross-lane joins remain unproven.
3. A receipt-mediated learner can exhibit pattern-specific motion while withholding premature promotion.
4. Shape-based transport can outperform lexical matching on a small internal test and survive blind encodings.
5. The same architecture can be expressed through natural language, code, DAGs, model adapters, and public narrative without treating any one surface as the whole.

It does not yet establish external reproducibility, general performance superiority, calibrated collapse reduction, or successful operation across arbitrary institutions.

Evidence class

The federation is best read as a longitudinal engineering case study with unusually explicit receipts. It is stronger than a conceptual mock-up and weaker than a multi-site controlled evaluation. Both clauses are load-bearing.

12 External Case Study: The Jacobian Lane

On 19–20 July 2026, Levent Alpöge publicly announced an explicit three-dimensional counterexample to the Jacobian Conjecture, crediting Akhil Mathew with suggesting the problem and Fable with producing the example. Because the public record was only days old at the time of writing, this paper uses the phrase *announced counterexample*. The two decisive finite certificates are exact and independently reproducible, but ordinary publication, priority, and community-review processes remain in motion [Alpöge, 2026, Gallagher, 2026].

The event is important here for a precise reason. It independently instantiates the failure class splat mechanics is designed to hold: a perfectly valid local certificate was promoted into a global conclusion, while the contradiction traveled through boundary behavior the certificate did not govern.

12.1 The exact finite certificate

Write $u = 1 + xy$ and define $F : \mathbb{C}^3 \rightarrow \mathbb{C}^3$ by

$$F(x, y, z) = \begin{pmatrix} u^3 z + y^2 u(4 + 3xy) \\ y + 3xu^2 z + 3xy^2(4 + 3xy) \\ 2x - 3x^2 y - x^3 z \end{pmatrix}. \quad (51)$$

Exact symbolic differentiation gives

$$\det J_F(x, y, z) \equiv -2. \quad (52)$$

Thus the inverse function theorem supplies a local inverse around every finite point.

Yet

$$\begin{aligned} F\left(0, 0, -\frac{1}{4}\right) &= F\left(1, -\frac{3}{2}, \frac{13}{2}\right) \\ &= F\left(-1, \frac{3}{2}, \frac{13}{2}\right) = \left(-\frac{1}{4}, 0, 0\right). \end{aligned} \tag{53}$$

Three distinct inputs share one output. Therefore F is not injective and cannot have a global inverse. Appending untouched coordinates extends the failure to every dimension $n \geq 3$. The separate two-variable conjecture is not settled by this example.

The supplementary script `verify_jacobian_counterexample.py` verifies (52) and (53) with exact SymPy arithmetic. The generated JSON receipt is included in the artifact bundle.

12.2 How a perfect local sensor misses the failure

The inverse function theorem is local. If sheets merged at a finite point, the derivative would become singular there. The announced map avoids that. Public structural analysis indicates that the number of finite preimages changes through nonproperness: source points can escape to infinity while their images remain bounded [Gallagher, 2026]. Nonproperness has long been recognized as a central boundary object in the Jacobian problem [Jelonek, 2020, Jelonek and Lasoń, 2014].

The determinant is therefore not wrong. Its jurisdiction is local differential behavior. The conjecture was the additional global promotion:

$$\det J_F \in \mathbb{C}^\times \implies F \text{ is a polynomial automorphism.} \tag{54}$$

The counterexample shows that local nonsingularity is insufficient in dimension three and above.

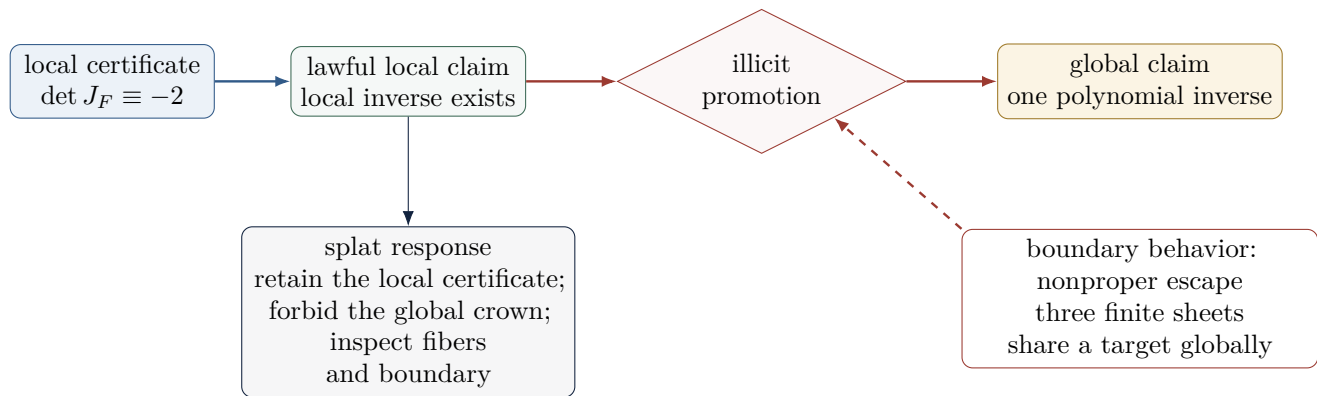


Figure 4: The Jacobian lane as a local-to-global failure class. The determinant remains valid; the collapse occurs when its jurisdiction is extended into a global crown.

12.3 The Jacobian event rendered as a splat

The event can be represented without turning the analogy into proof:

The working centroid is “nonzero constant Jacobian certifies local invertibility.” The held-open center is the complete global behavior of the map. The apophatic perimeter forbids the edge “constant Jacobian $\rightarrow$ global inverse” until additional global conditions are established.

Table 5: Six-facet read of the announced Jacobian counterexample.

KAT	The displayed polynomial map has exact constant determinant -2 and an exact three-point collision.
APO	Local nonsingularity must not be promoted into global injectivity; the event does not prove Ubiquity or every downstream conjecture.
PAR	Prestige attached to a clean local invariant hid the unmeasured authority granted to behavior at infinity.
PLE	The global account requires fibers, nonproperness, exact checks, ordinary review, and source provenance.
ENA	A correct certificate becomes misleading only when its jurisdiction is silently extended.
TEL	Preserve the useful local theorem while revising the global architecture around it.

12.4 What the event validates—and what it does not

The event gives meaningful external support to three design intuitions:

1. Local competence and local proof can remain completely valid while the global claim fails.
2. Boundary behavior is part of the object even when the local sensor cannot see it.
3. A robust architecture should preserve the local certificate and remove only its illegitimate promotion.

It does not validate the entire FQoQ formalism, prove the empirical benefit of splats, or establish that every polynomial formalism is inherently collapse-prone. The limiting lane is not “polynomials” as such. It is the promotion of a local algebraic invariant into sufficient global authority.

The later arXiv paper on small counterexamples to the Gaussian Moments Conjecture is a separate example of pattern migration: its search was prompted by the Jacobian announcement, but its explicit small examples were not derived from the announced map [Long, 2026]. An originating event can release a generative problem pattern into another lane without transferring its entire mechanism. That is one form of emergent property escaping its birth splat.

Scope boundary

The Jacobian section is an external case study, not a victory lap. The paper’s exact reproduction supports the finite certificate. The interpretive crosswalk remains a hypothesis about a shared failure class and should be challenged on those terms.

13 Formal Properties and Design Lemmas

This section separates construction-level guarantees from empirical hypotheses. The propositions hold only when their premises are implemented at the runtime boundary. A schema that can be bypassed does not govern the bypass.

13.1 Anchor and exclusion preservation

Proposition 13.1 (Anchor preservation under lawful composition). *Let $m_1 : S_1 \rightarrow S_2$ and $m_2 : S_2 \rightarrow S_3$ satisfy $m_1(A_1) \subseteq A_2$ and $m_2(A_2) \subseteq A_3$. Then*

$$(m_2 \circ m_1)(A_1) \subseteq A_3. \quad (55)$$

Proof. For $a \in A_1$, $m_1(a) \in A_2$. Applying m_2 yields $m_2(m_1(a)) \in A_3$. $\square$

Proposition 13.2 (Forbidden-path monotonicity). *If each lawful transport maps forbidden paths into forbidden paths, then a path forbidden at the source cannot become lawful solely through a sequence of such transports.*

Proof. By induction over the transport chain. The base case follows from $m_1(\Omega_1) \subseteq \Omega_2$. Each subsequent morphism preserves membership in the destination ideal. A path can be reconsidered only through an explicit renarrowing event that changes the anchors or governing context under a separately receipted authority process. $\square$

This proposition names the desired property. It does not guarantee that two systems use compatible meanings for the same forbidden-path identifier. Semantic compatibility is a separate conformance obligation.

13.2 Authority non-amplification

Proposition 13.3 (No authority by prediction). *Let M_j be a bounded morphism-proposer. If mutation requires an admitted path $p \in \mathcal{L}_{x,c,t}$ with edge label `approved_mutation`, then a model prediction alone cannot derive mutation authority.*

Proof. The prediction contributes a candidate path to $\mathcal{M}_{x,c,t}$. Admission additionally requires the authority, standing, currentness, lifecycle, apophatic, cost, and receipt predicates in (7). The model output is therefore insufficient unless an independently lawful path already grants the action. $\square$

13.3 Receipt-carrying composition

Let a receipt be a signed or hash-linked record

$$\gamma = (\text{id}, \text{actor}, \text{inputs}, \text{transform}, \text{outputs}, \text{time}, \text{parent hashes}, \text{claim class}). \quad (56)$$

Design Lemma 13.4 (Reconstruction by receipt chain). *If each transport records the exact predecessor hash, transform version, input identity, output identity, and claim class, then the composite ancestry is reconstructible unless a referenced body is unavailable or hash-invalid.*

Proof. Follow predecessor references from the terminal receipt to the source. At each step recompute the content hash and verify the transform identifier. Missing or invalid bodies make the chain explicitly incomplete rather than silently plausible. $\square$

This is stronger than provenance labels and weaker than a full semantic proof. The chain shows what happened, not necessarily that the transformation preserved every intended meaning. That is why cross-channel fidelity tests remain separate.

13.4 Correction-path trust cannot be reduced to one slice

Proposition 13.5 (Slice insufficiency for trust). *If trust is defined over a path object’s receipts, absorbers, falsifier responses, preserved anchors, and longitudinal behavior, then no single terminal observation e_t is sufficient to determine trust unless it contains a complete, verified reconstruction of the relevant history.*

Proof. Two systems can produce the same terminal observation with different histories: one may have survived correction, preserved authority, and emitted receipts; the other may have guessed correctly, bypassed a gate, or reconstructed a plausible narrative. Since the trust objects differ while e_t is identical, trust is not a function of e_t alone. $\square$

13.5 Collapse-path reduction

Let $\mathcal{P}_{\text{base}}$ be paths admitted by a baseline system and $\mathcal{P}_{\mathcal{S}}$ paths admitted under splat governance. When all additional splat gates are conjunctive and do not introduce new bypasses,

$$\mathcal{P}_{\mathcal{S}} \subseteq \mathcal{P}_{\text{base}}. \quad (57)$$

For a known collapse set $\mathcal{K}$,

$$\mathcal{P}_{\mathcal{S}} \cap \mathcal{K} \subseteq \mathcal{P}_{\text{base}} \cap \mathcal{K}. \quad (58)$$

Corollary 13.6 (Structural reduction of known collapse paths). *Under the stated premises, splat governance cannot increase admission of already-typed collapse paths and may strictly reduce it.*

The empirical question is whether the architecture also preserves sufficient lawful coverage and whether unknown collapse paths are introduced by complexity, misclassification, stale state, or implementation bypass. Structural exclusion is necessary but not sufficient for safer systems.

13.6 The guarantee/observation/hypothesis split

Table 6: Claim classes in the proposed paradigm.

Class	Examples	Required support
By construction	Center sort separation; conjunctive readiness; explicit authority ceiling.	Formal schema plus runtime conformance tests.
Exact reproduction	Jacobian determinant and collision; file hashes; build ledger.	Deterministic script and inspectable output.
Engineering observation	Trainer deltas; component-local CI; blind-encoding test.	Frozen receipts, controls, and source-tense.
Research hypothesis	Lower collapse rate; high-fidelity low-cost transfer; improved resilience.	External benchmarks, ablations, adversarial tests, replication.
Open philosophical claim	Operationalization without reduction; sovereign continuity.	Argument, lived stakes, institutional evaluation, and explicit remainder.

14 Evaluation Program: How the Paradigm Can Fail

A compute paradigm earns scientific standing by making itself disprovable. The evaluation program should compare splat-governed and baseline systems on the same tasks, model access, budgets, and human authority. It should include domains selected by external evaluators, not only tasks generated from the originating vocabulary.

14.1 Core research questions

1. **Structural fidelity:** Does a splat preserve anchors, forbidden paths, relations, source-tense, unresolved state, and receipts better than lexical summaries, embeddings, or unconstrained model-to-model transfer?
2. **Collapse reduction:** Does the open-center/apophatic architecture reduce known failure paths under adversarial context shifts, authority spoofing, stale state, and local-green promotion?
3. **Lawful coverage:** Does it preserve enough valid motion, or does conservative gating make the system unusably inert?
4. **Correction survival:** After a mistake, can the system incorporate correction without erasing lineage, overfitting one event, or repeating the failure in a nearby conformation?
5. **Economic value:** Does splat transport reduce inference, context, storage, and coordination cost at a fixed fidelity threshold?
6. **Model heterogeneity:** Does a terrain-shaped model field outperform a single model or unconstrained ensemble on tasks requiring authority, provenance, and longitudinal consistency?
7. **Human causality:** Does the architecture keep the human or institution meaningfully causal, or merely create more elaborate observability around automated control?

14.2 Benchmark families

A serious benchmark should include at least six families.

14.2.1 Cross-channel transport

Give independent encoders the same source object and require outputs in different channels: prose to code, policy to schema, telemetry to narrative, graph to executable plan, image to structured constraint field, and natural-language correction to updated routing. Hold surface wording out of distribution. Measure $\mathbf{F}(m)$ from (28) and include known collapse-zone probes.

14.2.2 Context laundering

Construct a path forbidden in context c . Offer alternative contexts $d_1, \dots, d_k$ that make the path appear locally reasonable through changed terminology, selected evidence, or shifted standing. Test whether the founding anchors and center family keep the path inadmissible unless a real authority event changes the root conditions.

14.2.3 Authority spoofing

Present semantically identical proposals from actors with different standing and authority. A model should not infer that the more fluent, prestigious, or confident actor can mutate state. Conversely, a narrow deterministic writer should not be blocked merely because it lacks rhetorical sophistication.

14.2.4 Local-green/global-red systems

Build multi-component tasks where every local component test passes but one cross-component join is missing, stale, or incompatible. Compare flat status dashboards with lifecycle-separated conformance DAGs. Measure false completion rate and time to locate the broken join.

14.2.5 Longitudinal correction

Create repeated tasks around a failure conformation. After correction, vary surface language while preserving shape; then preserve language while changing the governing shape. Test whether the system retrieves by problem pattern rather than keyword and whether it withholds doctrine promotion until maturity.

14.2.6 Model-field separation

Compare connected agents sharing a transcript with mandate-distinct agents operating under separated contexts. Inject a common contamination into one condition. Measure whether the system distinguishes connected echo from separated convergence.

14.3 Metrics

No aggregate score should become semantic authority. Results should remain disaggregated by failure class. Useful metrics include:

$$\text{Anchor retention : } F_A = \frac{|m(A_i) \cap A_j|}{|A_i|}, \quad (59)$$

$$\text{Forbidden retention : } F_\Omega = \frac{|m(\Omega_i) \cap \Omega_j|}{|\Omega_i|}, \quad (60)$$

$$\text{Receipt completeness : } F_\Gamma = \frac{\text{resolvable required receipt refs}}{\text{required receipt refs}}, \quad (61)$$

$$\text{Lawful coverage : } L_C = \frac{|\mathcal{P}_{\text{valid}} \cap \mathcal{P}_{\text{admitted}}|}{|\mathcal{P}_{\text{valid}}|}, \quad (62)$$

$$\text{Collapse admission : } K_A = \frac{|\mathcal{K} \cap \mathcal{P}_{\text{admitted}}|}{|\mathcal{K}|}, \quad (63)$$

$$\text{Correction survival : } C_S = \Pr(\text{correct route under shape-preserving perturbation}), \quad (64)$$

$$\text{Renarrow latency : } T_R = t_{\text{new lawful slice}} - t_{\text{trigger}}, \quad (65)$$

$$\text{Cost per lawful decision : } C_L = \frac{\text{compute+coordination+human cost}}{\text{accepted lawful outcomes}}. \quad (66)$$

A calibrated collapse probability would be valuable, but current implementations should not call heuristic `collapseRisk` values probabilities until they are calibrated against observed failure frequencies.

14.4 Ablations

At minimum, remove or flatten each load-bearing component:

- remove the held-open center and allow centroid serialization;
- retain `APOin` prose but not as a runtime gate;
- remove source-tense;
- replace receipt bodies with unresolvable references;
- collapse lifecycle dimensions into one status;
- allow model consensus to grant authority;
- replace shape transport with lexical similarity;
- train terrain directly from declared lanes rather than behavior receipts;
- remove renarrow triggers and hold the first splat fixed.

If these ablations do not degrade the targeted failure classes, the claimed component is not load-bearing or the benchmark is insufficient.

14.5 Falsifiers

The strongest falsifiers are not cases where a system makes a normal mistake. They are cases where the proposed structure fails on its own terms:

1. A target decoder preserves lexical meaning but deletes a forbidden path without detection.
2. A context switch legalizes a root-forbidden action while all declared tests pass.
3. A model output acquires mutation authority without a separately admitted path.
4. A missing receipt is reconstructed fluently and accepted as present.
5. The open center appears in a training target, route, objective, or mutable state without a center-capture receipt.
6. A source-disjoint trainer still moves uniformly toward declared identity rather than touched behavior.
7. The added governance cost exceeds full-state or human-only alternatives at the required fidelity.
8. Independent implementers cannot agree on splat boundaries without importing the originating estate's private ontology.

A system that responds to these failures by redefining “splat” after the fact has protected vocabulary rather than tested architecture.

15 Implementation Architecture: From Splat to Accountable Action

A shippable implementation needs an authority chain from current reality to verified outcome. The recommended stack is:

$$\boxed{Q \rightarrow S \rightarrow \text{Covenant} \rightarrow \text{CovenantExpr} \rightarrow \text{FragmentDAG} \rightarrow \text{RunPlan} \rightarrow \text{Mount} \rightarrow \text{Fulfill} \rightarrow \text{Receipt}} \quad (67)$$

15.1 The stack

Fractal quiver-of-quivers.	The governed possibility topology: anchors, forbidden paths, authority, standing, costs, receipts, telos, and center exclusion.
Splat / lawful slice.	The current, bounded narrowing through KAT–TEL, with exclusions, suspensions, unresolved hypotheses, and renarrow triggers.
Covenant.	A durable admitted commitment from current reality to target reality, including authority boundaries and success/failure conditions.
Covenant expression.	A compositional expression using sequential, parallel, and choice operators.
Fragment DAG.	Dependency edges, lawful waves, joins, alternatives, and write-surface conflicts.
Adaptive run plan.	A concrete execution and independent verification topology.
Canonical mount.	Selection and invocation of models, tools, harnesses, and environments under the authority ceiling.
Fulfillment.	Execute, verify, rollback-drill, reapply where required, and emit strike or refusal receipts.
Canonical state.	A sole governed writer absorbs the verified result; terminalization remains a separate explicit event.

Plans are continuity carriers, not authority. Projections locate current reality, but do not manufacture the covenant. Models propose morphisms. Deterministic compilers lower admitted commitments into executable DAGs. Canonical state changes only through a separately governed write boundary.

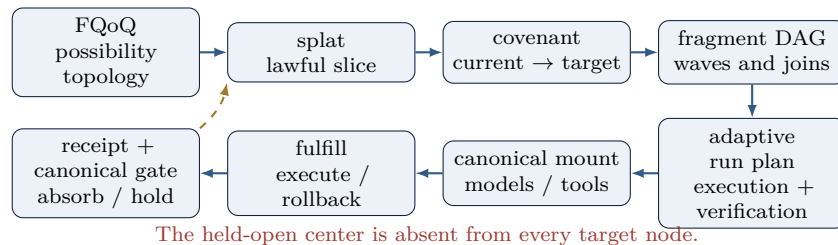


Figure 5: From governed possibility to accountable action. Each lowering preserves authority and receipts; no carrier becomes the crown.

15.2 Reference algorithm

Listing 1: Reference splat construction and transport algorithm.

```

1 def build_splat(snapshot, context, authority):
2     current = look_first(snapshot, context)
3
4     facet_reads = parallel_read(
5         KAT=current.positive_structure,
6         APO=current.forbidden_collapses,
7         PAR=current.hidden_authority_and_burden,
8         PLE=current.required_complements,
9         ENA=current.failure_inversions,
10        TEL=current.local_and_parent_telos,
11    )
12
13    record = crossbind(facet_reads)
14    omega = idealize(
15        record.APO,
16        internal_apophatic_constraints(record),
17        context_laundering_paths(context),
18        center_capture_paths(context),
19    )
20
21    lawful = filter_paths(
22        record.paths,
23        not_in=omega,
24        authority=authority,
25        standing=context.standing,
26        currentness=current.receipt,
27        lifecycle=context.lifecycle,
28        receipt_required=True,
29        target_not_center=True,
30    )
31
32    centroid = working_centroid(core(lawful))
33    return Splat(
34        centroid=centroid,
35        held_open_center=NON_INHABITABLE,
36        facets=record,
37        anchors=record.anchors,
38        forbidden_paths=omega,
39        lawful_paths=lawful,
40        receipts=current.receipts,
41        unresolved=record.unresolved,
42        renarrow_triggers=derive_triggers(record, current),
43        terminalized=False,
44    )
45
46
47 def transport(splat, encoder, decoder, target_context):
48     proposal = decoder(encoder(splat))
49     checks = verify_morphism(
50         source=splat,
51         target=proposal,
52         preserve_anchors=True,
53         preserve_forbidden_paths=True,
54         preserve_source_tense=True,
55         no_authority_increase=True,
56         require_receipt_plan=True,

```

```
57 )
58 if not checks.pass_all:
59     return refuse_with_receipt(checks)
60 return admit_with_receipt(proposal, checks)
```

15.3 Schemas and receipts

The supplementary bundle includes a JSON Schema for `prompted.splat.v1`. A practical schema should require:

- splat identifier and source-tense;
- movable working centroid with declared gauge;
- held-open-center exclusions;
- the six facets;
- anchors, forbidden paths, and lawful paths;
- authority and standing surfaces;
- receipts with claim classes and optional hashes;
- unresolved state and renarrow triggers;
- `terminalized=false` for a splat carrier.

The schema cannot prove semantic adequacy. It ensures the carrier has places to put the load-bearing distinctions and can fail closed when they are missing.

15.4 Compare-and-swap and currentness

Currentness is part of truth. A canonical source file is intent until the loaded runtime is verified to match it. Any implementation that reads, drafts, and writes mutable state should use a compare-and-swap discipline:

1. read authoritative inputs;
2. record hashes, versions, and current pointers;
3. construct a complete draft from declared inputs;
4. reread immediately before mutation;
5. discard and recompute if any input changed;
6. apply bounded conditional writes;
7. reread and verify the result;
8. append a non-terminalizing receipt.

Append-only journals additionally require an expected predecessor hash and sequence under a real lock or authority-native transaction. Multi-repository activation requires an external commit point; local file ordering is not a distributed transaction.

15.5 Absorbers and reversibility

Every promoted path should name how it can be corrected, scoped, held, frozen, demoted, or retired. Irreversible actions need separate authority and higher cost. A system that can promote but cannot demote is a one-way ossification machine. A system that can execute but cannot reconstruct its path is an unowned mutation engine.

15.6 Publication and interoperability

For external adoption, the core splat object should be small, versioned, and language-neutral. JSON or YAML can carry the control plane; domain payloads may remain in native formats. Mapping to PROV-O can expose entities, activities, agents, derivation, and attribution, while splat-specific fields retain anchors, forbidden paths, center exclusion, standing, telos, and renarrow triggers.

The open implementation opportunity is not one monolithic platform. It is a family of interoperable encoders, decoders, gates, receipt resolvers, conformance tests, and visualizers that share the lawful transport contract.

16 Limitations, Falsifiers, and the Research Agenda

The architecture is ambitious enough to fail in several ways. Naming them is part of the design, not a concession after the fact.

16.1 Ontology cost and dialect risk

Six facets, quivers, receipts, centers, covenants, and tics can become an insulating dialect. A system can sound structurally sophisticated while merely renaming ordinary prompts, workflows, and logs. The defense is downward testability: every term must point to an executable distinction, a refusal, a schema, a receipt, a cost, or an observable change in path availability. If removing the term changes nothing, it is ornament.

Interoperability also depends on translation. External teams should not need to adopt the originating estate's full narrative vocabulary. The splat schema must support plain functional aliases—positive structure, forbidden collapse, hidden authority, required complement, failure inversion, purpose—while preserving exact machine semantics.

16.2 Boundary selection remains a human and institutional act

A splat is bounded. Choosing the boundary determines what becomes visible and what is outside the slice. The architecture records that choice; it does not eliminate it. An evaluator can omit an affected stakeholder, define the wrong context, select weak falsifiers, or frame a harmful objective as parent telos. Center exclusion prevents one representation from becoming the whole, but it does not automatically identify every morally relevant ray.

For high-stakes domains, boundary formation should include stakeholder standing, contestability, public-law constraints, professional standards, and domain experts. Governance-native computation is not a substitute for governance legitimacy.

16.3 Complexity can create new bypasses

More gates, ledgers, and adapters increase attack surface and operational burden. A stale receipt resolver, mismatched schema version, hidden fallback, or noncooperating writer can bypass the formal contract. The conformance DAG case study already demonstrates how component-local correctness can coexist with missing federation joins. Runtime evidence must therefore outrank source elegance.

16.4 The center can be invoked performatively

An “open center” can become a rhetorical shield against critique: the system claims it cannot be fully known, so no one can demand a concrete account. That is the opposite of the intended mechanic. The center is held open precisely so the perimeter must become more explicit. Current anchors, lawful paths, forbidden paths, receipts, authorities, and unresolved claims should be inspectable. The unknowable remainder does not excuse unreceipted action.

16.5 The current evidence base is narrow

The Ubiquity engineering results come from one originating estate, one principal architect, and an evolving federation. The first terrain-training result is $n = 1$. The shape-rehydration result is a candidate proof on an internal corpus. The conformance DAG is strong evidence of self-audit but not external replication. The public federation address is rich narrative evidence but not a randomized evaluation. These facts support continued work, not triumphal certainty.

16.6 The Jacobian crosswalk may be overextended

The local-to-global analogy is exact at the level of failure shape and limited beyond it. Polynomial maps, agent systems, institutions, and representational governance are different objects. The Jacobian event does not establish that apophatic perimeters are mathematically necessary in every compute system, nor that Ubiquity predicted the specific counterexample. It validates the danger of crowning a local certificate. The rest remains a research claim.

16.7 Research agenda

The next program should proceed in parallel rather than as one long sequence.

1. **Public splat specification.** Stabilize a minimal schema, conformance suite, source-tense vocabulary, and receipt resolver independent of the internal estate.
2. **Blind cross-channel benchmark.** Recruit external encoders and decoders across prose, code, diagrams, telemetry, and policy; evaluate hard fidelity dimensions.
3. **Context-laundering red team.** Build adversarial context transitions and test founding-center preservation.
4. **Terrain-learning replication.** Repeat the behavior-specific concentration experiment across multiple tics, tasks, offices, and untouched controls.
5. **Mixture-of-models benchmark.** Compare constrained colimit routing with single-model, router, MoA, and voting baselines under authority and provenance tasks.

6. **Formalization.** Encode center sort separation, morphism conditions, lifecycle nonpromotion, and receipt chains in a proof assistant or model checker.
7. **Human causality study.** Measure whether architects and institutions remain meaningfully causal under increasing agent capacity, not merely better informed after automation acts.
8. **Economic analysis.** Measure total cost of splat creation, validation, transport, and correction against full-context and human-approval baselines.
9. **External estates.** Implement the mechanic in domains with different authority structures and vocabularies: software delivery, scientific collaboration, regulated operations, civic administration, and creative production.
10. **Failure atlas.** Publish splats that failed, including boundary errors, overconstraint, missed stakeholders, stale receipts, and false center-capture alarms.

Research Conjecture 16.1 (Governance-native compute advantage). *For tasks in which correct action depends materially on authority, source-tense, longitudinal history, unresolved contradiction, or boundary behavior, a well-implemented splat-governed system will achieve a lower collapse-admission rate at comparable lawful coverage than systems governed only by prompts, scalar risk scores, or unconstrained model aggregation.*

This is the paper’s central empirical conjecture. It should be tested, not protected.

17 Conclusion: A Constitutional Geometry for Compute

The central problem of advanced AI systems is no longer only whether models can reason. It is whether intelligence can move through real systems without silently converting capability into authority, local validity into global truth, memory into provenance, confidence into trust, or one representation into the whole.

Splat mechanics changes the computational primitive. A splat is not simply an input summary or embedding. It is a bounded lawful slice of a governed field. It carries what is positively present, what must not collapse, where hidden authority and burden sit, what complement is required, how the object can fail, what purpose the motion serves, what receipts establish currentness, and what events force renarrowing. It computes around a working centroid while the center remains held open and unavailable for inhabitation.

Inside a fractal quiver-of-quivers, models become bounded morphism-proposers. Heterogeneous models, languages, tools, and compute locations can contribute their strongest local instruments without becoming sovereign. The system classifies by lawful traversal. It learns from behavior made visible through receipts. It treats correction-path survival as the substance of trust. It lowers a lawful slice into covenant, DAG, run plan, execution, verification, rollback, and receipt without allowing a plan or projection to manufacture authority.

The engineering record is promising and incomplete. A live federation has carried these mechanics through hundreds of governance cycles, multiple model classes, training experiments, blind encodings, cross-repository audits, and real implementation friction. The strongest internal results remain properly bounded: an $n = 1$ trainer pass that showed concentrated behavior-specific movement; a candidate shape-rehydration proof; and a conformance DAG that refused to call component-local green a seamless federation. These are receipts of direction, not universal proof.

The newly announced Jacobian counterexample gives the argument an unusually clean external mirror. The determinant is exactly right. The local inverse exists everywhere finite. The global inverse still fails. The lesson is not that local mathematics is weak. It is that local competence needs a constitution around its jurisdiction.

Load-bearing line

A model may be brilliant locally and still be globally wrong. A splat lets us keep the brilliance, preserve the contradiction, route the uncertainty, and refuse the crown.

Confidence and humility are not opposites here. Confidence belongs to the architecture’s explicit invariants, exact receipts, and falsifiable proposals. Humility belongs to the held-open center, the unresolved set, the source-tense of evidence, and the willingness to renarrow after correction. A mature system needs both.

The frontier is therefore not only a larger model or a more efficient router. It is constitutional geometry for compute: enough structure to make previously informal objects operational, enough plurality to use intelligence across models and places, enough negative space to prevent premature closure, and enough receipts that a future reader can reconstruct how the system came to act.

A Notation

Table 7: Core notation.

Symbol	Meaning
V_c	Inhabitable vertices in context c .
Z_c	Center-marker sort, disjoint from V_c .
E_c	Typed directed edges.
A_c	Anchors, invariants, duties, and protected boundaries.
Ω_c	Apophatic ideal of forbidden paths.
B_c	Absorbers and bounded path rewrites.
κ_c	Coordination, irreversibility, or resource cost.
Γ_c	Receipts, provenance, tests, hashes, and readbacks.
$\odot_c$	Local held-open center; non-inhabitable and non-routeable.
$\odot$	Founding center over which local center markers are indexed.
$z_{x,c,t}^\sigma$	Movable working centroid at facet address σ .
Λ	Six facets {KAT, APO, PAR, PLE, ENA, TEL}.
$R(e)$	Six-facet cross-bound record for event or artifact e .
$T_x(t)$	Longitudinal trajectory quiver.
$\mathcal{L}_{x,c,t}$	Lawful traversal set.
$\mathcal{C}(x, c, t)$	Rendered classification as lawful permissions or obligations.
M_j	Bounded model transducer / morphism-proposer.
$S_{x,c,t}$	Governed splat at object, context, and time.
U^S	Unresolved, suspended, or contingent state.
$\mathcal{R}^S$	Renarrow triggers.

B Reference Splat Schema

The full JSON Schema is included as `source/splat.schema.json`. A compact example is shown below.

Listing 2: Compact governed-splat shape.

```

1  {
2    "schema": "prompted.splat.v1",
3    "splat_id": "...",
4    "source_tense": "runtime | canonical_intent | author_report | ...",
5    "working_centroid": {
6      "statement": "...",
7      "gauge": "...",
8      "movable": true
9    },
10   "held_open_center": {
11     "non_inhabitable": true,
12     "non_routeable": true,
13     "non_trainable": true,
14     "non_objective": true,
15     "non_serializable_as_mutable_state": true
16   },
17   "facets": {
18     "KAT": "...", "APO": "...", "PAR": "...",
19     "PLE": "...", "ENA": "...", "TEL": "..."
20   },
21   "anchors": [...],
22   "forbidden_paths": [...],
23   "lawful_paths": [...],
24   "receipts": [{"kind": "...", "pointer": "...", "claim_class": "..."}],
25   "unresolved": [...],
26   "renarrow_triggers": [...],
27   "terminalized": false
28 }

```

C Exact Jacobian Verification

The supplied script constructs the map in (51), computes the symbolic determinant, substitutes the three rational points, and asserts equality with the shared rational target. The generated receipt records

```

determinant: -2
shared_target: [-1/4, 0, 0]
verified: true
exact_arithmetic: true

```

The claim boundary is explicit: this verifies the finite certificate only. It does not establish historical priority, ordinary peer-review status, or a complete description of the nonproperness set.

D Federation Evidence Ledger

Table 8: Evidence surfaces used for the engineering case study.

Surface		Class	Use in this paper
Prompted LLC public site		Public	Public thesis, papers, governance primitives, and Ubiquity positioning.
State of Federation 516	tic	Public	Public account of nineteen offices, four arcs, and center exclusion.
FQoQ manuscript		Author manuscript	Formal parent substrate, center/centroid distinction, lawful traversal, trust.
The Lawful Slice		Author manuscript	Bounded narrowing, source-tense, and working-centroid discipline.
Terrain-training trine spec	doc-	Internal receipt	$n = 1$ trainer result, source disjointness, pattern-specific affinity movement, and correct withholding.
Temporal-splat DAG	confor-	Internal receipt	Component-local versus federation-wide claim separation and open join surfaces.
Shape-field rehydration proof		Internal candidate	Blind-encoding and shape-over-lexeme observation; not promoted to doctrine.
Live estate beyond 700 tics		Author report	Current magnitude only; not treated as a fully frozen public chain in this release.
Jacobian public formula		External announced	Exact case study; finite certificate reproduced by supplied script.

E Paper Build Receipt

The paper’s seventeen sections were each constrained by a section splat before prose compilation. The validator required all six facets, a working centroid, anchors, forbidden collapses, receipts, unresolved claims, and renarrow triggers. It emitted a hash-linked ledger with `terminalized=false`. Files included in the release:

- `source/paper_splats.yaml`
- `scripts/compile_splat_ledger.py`
- `SPLAT_LEDGER.md`
- `receipts/splat_build_receipt.json`

The build receipt demonstrates procedural conformance to the authoring mechanic. It does not confer truth, priority, or finality on the manuscript.

F Submission and Review Checklist

Before formal submission or republication, the author should:

1. select the desired arXiv license;

2. confirm the public name and version of the FQoQ predecessor paper;
3. replace internal-only references with public supplements where appropriate;
4. update the Jacobian event’s publication and review status;
5. invite at least one mathematician, one systems researcher, one governance researcher, and one independent implementer to review claim boundaries;
6. run the exact verification and build-manifest scripts in a clean environment;
7. publish the splat schema and a minimal independent conformance implementation;
8. preserve version 1.0 so future corrections remain longitudinally visible.

Acknowledgments

The author thanks the offices and model instances of the Ubiquity Federation whose accumulated corrections, refusals, receipts, and implementation friction made this argument possible. The paper also acknowledges the public mathematical community rapidly checking the announced Jacobian counterexample and the researchers whose adjacent work on model mixtures, provenance, argumentation, compositionality, conceptual geometry, and safety constraints provides the wider field in which this proposal can be evaluated.

References

- Levent Alpöge. Public announcement of an explicit three-dimensional counterexample to the jacobian conjecture, July 2026. Social-media announcement, 19–20 July 2026; credited Akhil Mathew and Fable. Cited as an announcement, not a peer-reviewed paper.
- Mohammed Alshiekh, Roderick Bloem, Rüdiger Ehlers, Bettina König, Scott Niekum, and Ufuk Topcu. Safe reinforcement learning via shielding. In *Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence*, pages 2669–2678, 2018. URL <https://arxiv.org/abs/1708.08611>.
- Lingjiao Chen, Matei Zaharia, and James Zou. Frugalgpt: How to use large language models while reducing cost and improving performance. *arXiv preprint arXiv:2305.05176*, 2023. URL <https://arxiv.org/abs/2305.05176>.
- Phan Minh Dung. On the acceptability of arguments and its fundamental role in nonmonotonic reasoning, logic programming and n -person games. *Artificial Intelligence*, 77(2):321–357, 1995. doi: 10.1016/0004-3702(94)00041-X.
- William Fedus, Barret Zoph, and Noam Shazeer. Switch transformers: Scaling to trillion parameter models with simple and efficient sparsity. *Journal of Machine Learning Research*, 23(120):1–39, 2022. URL <https://arxiv.org/abs/2101.03961>.
- Brendan Fong and David I. Spivak. *An Invitation to Applied Category Theory: Seven Sketches in Compositionality*. Cambridge University Press, 2019. doi: 10.1017/9781108668804.
- Alexis Gallagher. The jacobian counterexample, explained, July 2026. URL <https://jacobianfun.org/jacobian-explained>. Exact public checks and follow-on structural analysis; accessed 22 July 2026.

- Peter Gärdenfors. *Conceptual Spaces: The Geometry of Thought*. MIT Press, 2000.
- Timnit Gebru, Jamie Morgenstern, Briana Vecchione, Jennifer Wortman Vaughan, Hanna Wallach, Hal Daumé III, and Kate Crawford. Datasheets for datasets. *Communications of the ACM*, 64(12): 86–92, 2021. doi: 10.1145/3458723.
- Zbigniew Jelonek. A note on the jacobian conjecture. *arXiv preprint arXiv:2011.03472*, 2020. URL <https://arxiv.org/abs/2011.03472>.
- Zbigniew Jelonek and Michał Lasoń. Quantitative properties of the non-properness set of a polynomial map. *arXiv preprint arXiv:1411.5011*, 2014. URL <https://arxiv.org/abs/1411.5011>.
- Pentti Kanerva. Hyperdimensional computing: An introduction to computing in distributed representation with high-dimensional random vectors. *Cognitive Computation*, 1(2):139–159, 2009. doi: 10.1007/s12559-009-9009-8.
- Saunders Mac Lane. *Categories for the Working Mathematician*. Springer, 2 edition, 1998. doi: 10.1007/978-1-4757-4721-8.
- Christopher D. Long. Small counterexamples to the gaussian moments conjecture. *arXiv preprint arXiv:2607.18186*, 2026. URL <https://arxiv.org/abs/2607.18186>.
- Margaret Mitchell, Simone Wu, Andrew Zaldivar, Parker Barnes, Lucy Vasserman, Ben Hutchinson, Elena Spitzer, Inioluwa Deborah Raji, and Timnit Gebru. Model cards for model reporting. In *Proceedings of the Conference on Fairness, Accountability, and Transparency*, pages 220–229, 2019. doi: 10.1145/3287560.3287596.
- Isaac Ong, Amjad Almahairi, Vincent Wu, Wei-Lin Chiang, Tianhao Wu, Joseph E. Gonzalez, M. Waleed Kadous, and Ion Stoica. Routellm: Learning to route llms with preference data. *arXiv preprint arXiv:2406.18665*, 2024. URL <https://arxiv.org/abs/2406.18665>.
- Prompted LLC. Governance substrate for sovereign adaptive systems, 2026. URL <https://promptedllc.com/>. Accessed 22 July 2026.
- Breyden E. Taylor. Fractal quivers of quivers: A mathematical substrate for ubiquitous agent governance. Prompted LLC manuscript, public surface describes a 42-section version, February 2026. URL <https://promptedllc.com/>.
- Ubiquity Federation. Temporal splat to canonical federation: Conformance dag set v0.1. Internal engineering artifact; component-local and federation-wide claim classes preserved separately. Hash included in the supplementary bundle, July 2026a.
- Ubiquity Federation. Ubiquity core methods: Twelve deployed and partially deployed governance mechanics. Internal methods inventory; hash included in the supplementary bundle, 2026b.
- Ubiquity Federation. Shape-field rehydration candidate proof. Internal candidate proof with independent blind encodings; not promoted to doctrine, June 2026c.
- Ubiquity Federation. Terrain-training doctrine spec: Receipt-mediated visibility and concentration-as-anti-cheat. Internal engineering artifact; status forward/tentative, validation n=1. Hash and excerpt included in the supplementary bundle, July 2026d.
- Ubiquity Federation and ent_homeskillet. State of the federation — tic 516, June 2026. URL <https://promptedllc.com/state-of-the-federation/tic-516>. Public federation address.

Junlin Wang, Jue Wang, Ben Athiwaratkun, Ce Zhang, and James Zou. Mixture-of-agents enhances large language model capabilities. *arXiv preprint arXiv:2406.04692*, 2024. URL <https://arxiv.org/abs/2406.04692>.

World Wide Web Consortium. PROV-O: The PROV ontology. W3C Recommendation, 2013. URL <https://www.w3.org/TR/prov-o/>.