

WORKING PAPER

September 2026

LiO v0.1: Deterministic Permission Composition Across Indispensable Inferential Bridges

Micheal Charles Preble

Independent Researcher

LiO Research Program · Version 0.1 · Public Scholarly Circulation

Prepared for SSRN submission and external scholarly review.

Table of Contents

Abstract.....3

1. The Inference-Control Problem.....3

2. The Minimal Primitive.....5

3. Validation Evidence.....7

4. Upstream Preconditions.....9

5. Limits and Non-Claims.....12

6. Residual Research Questions.....15

7. Discussion and Conclusion.....18

Appendix A. Claim Ledger.....19

References.....20

Evidence-control note.....21

Author note

This independent working paper reflects the views of the author. It is released for scholarly review. It does not claim mathematical novelty, patentability, universal correctness, domain-specific validity, or independent human validation. The composition kernel and E12 Cycle 1 result are reported within the scope stated in the paper.

Abstract

This paper introduces LiO v0.1, a minimal computational primitive for constraining downstream assertions in multi-stage derivations. Given a declared derivation with an explicitly supplied nonempty set of indispensable inferential bridges, each resolved bridge is assigned an upstream-adjudicated permission set over a finite assertion universe. LiO composes those permissions by intersection and preserves unresolved bridge state explicitly. At the target level, the mechanism reduces to coordinatewise three-valued conjunction: an assertion class is unauthorized if any resolved indispensable bridge excludes it, authorized if every indispensable bridge is resolved and permits it, and unresolved otherwise.

In E12 Review Cycle 1, a separate model reimplemented LiO from the frozen specification without access to the creator reference implementation, held-out expected outputs, creator predictions, or creator-side break ledger. The cross-model implementation reproduced all seven public cases exactly on the fields it implemented. After implementation freeze, it then matched creator-frozen held-out results on all twelve H001–H012 cases for `exact_final_permissions` and `target_authorized`. These results support cross-implementation reproducibility of the tested permission-composition behavior. They do not establish full output-contract conformance, domain-specific validity, universal correctness, independent human validation, or novelty of the underlying set-theoretic operation.

LiO v0.1 does not determine which semantic distinctions belong in the assertion universe, whether the declared derivation and indispensable dependency set are correct, or how bridge permissions should be adjudicated from evidence. Its present contribution is narrower: given a declared admissible derivation, a complete supplied set of indispensable bridges, and upstream-adjudicated permissions over a common assertion universe, LiO deterministically computes the downstream authorization state that survives every indispensable bridge. This separation yields a larger research program organized around representation, adjudication, and composition.

1. The Inference-Control Problem

1.1 The gap between derivation and authorization

Scientific and technical arguments routinely depend on chains of inferential transitions. A premise is interpreted, mapped, transformed, generalized, translated, or otherwise carried across one or more bridges before a downstream claim is stated. The fact that a derivation can be written as a sequence does not by itself establish that every class of downstream assertion remains authorized at the end of that sequence. A later claim may be stronger in scope, causality, modality, evidential status, population, or conditionality than one or more indispensable steps legitimately permit.

The problem addressed here is therefore not whether a conclusion is globally true, nor whether a body of evidence is sufficient under a domain-specific scientific standard. The narrower problem is inference control: once a derivation has been declared and the permissions of its indispensable inferential bridges have been adjudicated upstream, how should those permissions be composed so that the downstream assertion cannot exceed what every indispensable bridge permits?

Different research traditions track dependencies and belief revision [1,2], abstract approximation and lattice structure [3], provenance [5], semantic correspondence across heterogeneous representations [6,7], confidence in structured arguments [12], and authorization or authority attenuation [8,9]. LiO v0.1 isolates a smaller question from that larger landscape. It asks what remains authorized after all indispensable bridge permissions have been taken into account, without allowing support elsewhere in the derivation to compensate for an explicit restriction imposed by one of those bridges.

1.2 Why non-compensatory composition matters

For assertion authorization, the **weakest-link principle** means that the authorization scope of a derivation cannot exceed what every indispensable bridge permits; support elsewhere cannot compensate for an explicit restriction imposed by another indispensable bridge.

Consider a declared derivation with two indispensable bridges. Let the assertion universe contain three classes: descriptive, associational, and causal. Suppose bridge b_1 permits descriptive and associational, while bridge b_2 permits descriptive and causal. Their jointly surviving permission space is

$$P(b_1) \cap P(b_2) = \{\text{descriptive}\}.$$

A target assertion of class causal is therefore unauthorized under that derivation. Bridge b_1 excludes the causal class, and no permission supplied by b_2 can restore what b_1 withholds. This is a non-compensatory rule: the composition result is constrained by every indispensable bridge rather than by an average, score, vote, confidence total, or prestige-weighted aggregation.

1.3 Representation, adjudication, and composition

The architecture separates naturally into three problems. **Representation** determines what assertion classes and semantic distinctions are present in the universe A . **Adjudication** determines what each bridge permits, excludes, or leaves unresolved. **Composition** determines what authorization survives once those bridge-local judgments have been supplied.

LiO v0.1 addresses the third problem. The first two are not treated as solved components hidden inside the primitive. They are explicit upstream dependencies whose adequacy determines whether deterministic composition is scientifically useful. This separation is central to the paper because it allows the composition mechanism to remain small while making the larger research problem more precise.

1.4 Related work and claim boundary

LiO is adjacent to several mature research traditions, and the present paper treats those traditions subtractively rather than using their existence to imply novelty by recombination. Truth-maintenance and belief-revision systems established long ago that reasoning systems can record reasons or dependencies and revise belief states when assumptions or supporting propositions change [1,2]. Those systems are important prior art for dependency tracking, reason maintenance, and revision. LiO therefore does not claim novelty for recording dependencies, propagating revisions, or conditioning downstream state on upstream reasons.

The mathematical kernel also sits inside familiar territory. Abstract interpretation formalized program-analysis frameworks over lattices of approximate assertions and studied monotone, compositional operations on those abstractions [3]. Strong three-valued logics, including Kleene-style systems, provide longstanding semantics for reasoning with a third state such as unknown [4]. These traditions undercut any novelty claim based on set intersection, lattice meet, non-expansion, monotonic restriction, or three-valued conjunction considered in isolation. LiO's current formal contribution is therefore not presented as a new algebra. Algebraic subsumption constrains claims about operator novelty, but it does not by itself determine whether a domain-specific interface, explicit precondition decomposition, or testable reference instantiation has independent scientific or practical value.

Provenance and ontology-matching research address neighboring but distinct problems. Database provenance characterizes where derived data came from and which source data contributed to its existence [5]. Ontology matching identifies semantic correspondences across heterogeneous ontologies, while alignment-repair work addresses incoherent mappings between them [6,7]. LiO does not replace these functions. It assumes that participating bridge permissions have already been mapped into a common assertion universe and uses provenance as a runtime obligation rather than as the composition rule itself.

A second neighboring cluster concerns evidence, warrants, and structured assurance. Research on evidence claims asks whether a claim is justified by the totality and quality of the evidence [10], and work on argument structure emphasizes that empirical conclusions depend on explicit inferential, theoretical, procedural, and contextual warrants [11]. Contemporary assurance research has also developed compositional methods for propagating confidence through structured claims and evidence while preserving warrants, context, and provenance [12]. Most directly, Li's 2026 preprint frames scientific claim calibration in terms of evidence licensing some forms of assertion while withholding others [13]. The general proposition that evidence constrains what may responsibly be claimed is therefore not a LiO novelty claim.

Authorization, policy-combination, and programming-language research provide additional applied precedents for permission propagation and non-compensatory constraint. Flow-Limited Authorization integrates authorization with information-flow control and studies delegation and revocation of authority [8]. Effect systems can track bounds on what computations or modules are permitted to do, including explicit reasoning about authority and attenuation [9]. XACML 3.0 defines policy- and rule-combining algorithms including deny-overrides, under which a Deny decision takes priority, while preserving richer extended Indeterminate{P}, Indeterminate{D}, and Indeterminate{DP} states [16]. Ramli, Nielson, and Nielson formally characterize XACML 3.0 component evaluation and combining operators, including an equivalent lattice-based characterization [17]. These are close precedents for non-compensatory decision composition with explicit indeterminacy, but they are not identical to LiO's flat TRUE/FALSE/UNRESOLVED target semantics or its derivation-relative assertion-class permission object. WS-Policy 1.5 defines a commutative policy-intersection operation that retains mutually compatible policy alternatives [18]; the W3C specification explicitly notes that this operation is only analogous at times to set intersection and does not imply formal set-intersection semantics. It is therefore a structural predecessor rather than a demonstrated identity. LiO does not claim novelty for permission sets, deny/permit combining, policy intersection, authorization logic, least privilege, attenuation, or propagation of constraints as general mechanisms.

The composition operator is not presented as a novel authorization or policy-combining algorithm. The narrower object studied here is a constrained reference instantiation in which each supplied indispensable inferential bridge is associated with an upstream-adjudicated permission set over explicit downstream assertion classes, those permissions are composed non-compensatorily, and unresolved bridge state is preserved. The literature above establishes substantial overlap in both ingredients and applied patterns. This review does **not** establish that no exact predecessor exists, and the present paper makes no novelty or patentability claim from the absence of one in this bounded citation pass. Its empirical claim remains behavioral: a separate model reconstructed the frozen composition specification and reproduced the Cycle 1 corpus on the implemented core fields.

2. The Minimal Primitive

2.1 Formal definition

Let D be a declared derivation. Let A be a finite assertion universe, and let $I_D \neq \emptyset$ be the supplied set of bridges declared indispensable to D . Each bridge $b \in I_D$ is either **RESOLVED**, with a permission set

$$P(b) \subseteq A,$$

or **UNRESOLVED**, in which case no exact permission set is supplied for that bridge.

For a resolved bridge, membership $a \in P(b)$ means that upstream adjudication of that inferential transition positively permits assertion class a to pass that transition. It does not mean merely that the bridge failed to object to a , and it does not mean that the bridge independently proves a . LiO consumes this local permission judgment; it does not generate it.

The primitive is defined relative to a declared derivation. It does not infer indispensability from natural language, discover omitted dependencies, adjudicate scientific truth, or determine whether the derivation itself is valid. Base assertions or axioms are outside this primitive and must be represented separately.

2.2 Fully resolved composition

If every bridge in I_D is resolved, LiO computes the exact final permission set

$$P_{final}(D) = \bigcap_{b \in I_D} P(b).$$

For a target assertion class $a \in A$,

$$Auth(a \mid D) = \begin{cases} \text{TRUE}, & a \in P_{final}(D), \\ \text{FALSE}, & a \notin P_{final}(D). \end{cases}$$

For example, let

$$A = \{\text{descriptive}, \text{causal}\},$$

$$P(b_1) = \{\text{descriptive}\}, P(b_2) = \{\text{descriptive}, \text{causal}\}.$$

Then

$$P_{final} = \{\text{descriptive}\},$$

and a causal target is unauthorized.

2.3 Partially unresolved composition

Let $R \subseteq I_D$ be the subset of resolved indispensable bridges. Let U denote the unresolved subset obtained by removing R from I_D . Define the known upper bound

$$K(R) = \begin{cases} A, & R = \emptyset, \\ \bigcap_{b \in R} P(b), & R \neq \emptyset. \end{cases}$$

Every completion of the unresolved bridge permissions must lie within $K(R)$. Thus $K(R)$ is the current may-permission space under the declared derivation.

For a target assertion class a ,

$$Auth(a \mid D) = \begin{cases} \text{FALSE}, & a \notin K(R), \\ \text{TRUE}, & U = \emptyset \text{ and } a \in K(R), \\ \text{UNRESOLVED}, & a \in K(R) \text{ and } U \neq \emptyset. \end{cases}$$

A resolved bridge can therefore determine exclusion even while other indispensable bridges remain unresolved. Conversely, unresolvedness is not silently treated as permission or rejection. If $K(R) = \emptyset$, the final permission set is already known to be empty even if some bridges remain unresolved, because no later resolution can restore an assertion excluded by the resolved intersection.

2.4 Coordinatewise three-valued conjunction

For each assertion class $a \in A$, define the bridge-level state

$$v_b(a) = \begin{cases} 1, & b \text{ is resolved and } a \in P(b), \\ 0, & b \text{ is resolved and } a \notin P(b), \\ ?, & b \text{ is unresolved.} \end{cases}$$

LiO v0.1 then reduces at the target level to

$$\text{Auth}(a \mid D) = \bigwedge_{b \in I_D} v_b(a).$$

The conjunction uses three states: 0 dominates, all 1 values yield 1, and any remaining case yields ?. In other words, one resolved exclusion is sufficient for determinate non-authorization; full authorization requires every indispensable bridge to be resolved and permitting; and otherwise the result remains unresolved.

At the universe level, the resolved mechanism is equivalent to coordinatewise conjunction over the indicator vectors of the permission sets. The underlying set-theoretic operation is therefore simple and familiar. Lattice-based approximation frameworks and strong three-valued logics substantially predate LiO [3,4], while bounded effect systems and access-control policy-combining algorithms provide applied precedents for attenuation and non-compensatory decision composition [9,16,17]. LiO does not claim mathematical novelty or novelty of the underlying combining operator. Algebraic subsumption settles important operator-level novelty questions; it does not by itself answer whether the derivation-specific interface, explicit upstream preconditions, and testable reference instantiation are useful.

2.5 Mapping to frozen invariants

The following are **frozen specification properties**, not a claim that E12 Cycle 1 exercised every property through the cross-model review. The frozen specification includes algebraic, semantic, metamorphic, provenance, revision, and conformance obligations. The composition/core behavior includes the following properties: non-expansion under intersection (I02), exact intersection for resolved indispensable bridges (I03), commutativity (I04), associativity (I05), idempotence (I06), monotonicity under permission restriction (I07), neutral full permission (I08), empty-set annihilation (I09), exclusion of non-indispensable bridges from the composition (I10), preservation of unresolvedness (I11), determinate exclusion under partial information (I12), non-negation (I13), confidence non-authority (I14), domain-label invariance under bijective renaming (I15), and absence of an inferred hidden total order among arbitrary assertion symbols (I20).

Canonical determinism, representation/canonicalization requirements, provenance completeness, revision-triggered recomputation, and fail-closed malformed-input behavior are additional frozen obligations associated with I01 and I16–I19. These requirements matter to a conforming implementation and runtime but should not be confused with the set-intersection kernel itself.

2.6 What the primitive is and is not

LiO v0.1 is a deterministic, non-compensatory composition mechanism over explicit bridge permission states. It enforces a derivation-relative constraint: the downstream assertion space cannot exceed the permissions supplied by every indispensable bridge in that derivation, and explicit unresolvedness is preserved rather than converted into approval or rejection.

LiO v0.1 is not a method for defining the correct assertion universe, adjudicating bridge permissions from evidence, validating the completeness of the indispensable dependency set, proving the admissibility of the derivation, evaluating alternative derivations, or determining global scientific truth. Those are separate problems.

3. Validation Evidence

3.1 E12 Cycle 1: blind cross-model reimplementation

E12 Review Cycle 1 tested whether the frozen LiO v0.1 composition behavior could be reconstructed by a separate model from the specification. The reviewer was Mistral Medium 3.5 in the Vibe Work environment. The reviewer received the frozen specification and seven designated public cases but did not receive the

creator reference implementation, creator-held expected outputs, creator predictions, novelty materials, or the creator-side failure/break ledger.

Before implementing the public cases, the reviewer recorded eleven specification ambiguities and froze its chosen interpretations. It then implemented the mechanism in the separate model environment and reproduced all seven public cases exactly on the fields it elected to implement: `exact_final_permissions` and `target_authorized`.

Before the held-out corpus was released, the reviewer reported a 64-character SHA-256 digest for its implementation:

```
0c4b941bfb5659ee9bb7f7d2b2037de65ff7edbac72e09c5d27b30e68ecb3883
```

The digest format was verified, but the creator did not independently hash the exact implementation bytes. It should therefore be treated as reviewer-reported implementation-freeze metadata rather than creator-verified custody of the implementation artifact.

3.2 Held-out blind execution

After implementation freeze, the reviewer received twelve held-out inputs, H001–H012, without creator expected outputs. The creator-side held-out input artifact had previously been frozen with SHA-256:

```
6df521895f2d6b6d7a67328a4fe2016dc3f4db358bfcbb6cf4fe7218962edca5
```

The reviewer executed its frozen implementation and returned outputs for all twelve cases. Comparison against the creator-frozen expectations produced 12/12 exact matches on the implemented core fields `exact_final_permissions` and `target_authorized`.

The raw output artifact actually received by the creator was independently measured at 932 bytes with SHA-256:

```
274c6ef01db8179ac5b52d6f17a575fe9a4787f5e9c73802f3dcfa6622a595d6
```

Reviewer-reported hash metadata for that raw output file was inconsistent with the bytes actually received and was preserved as a process discrepancy rather than silently reconciled. The discrepancy did not change the semantic comparison: the received twelve output records matched the creator-frozen expected core fields exactly.

3.3 What Cycle 1 tested

The E12 Cycle 1 cross-model corpus contained nineteen cases: seven public and twelve held-out. This number describes only the cross-model review corpus and not the larger creator-side validation suite, which includes additional fixed-gold, property-based, malformed-input, domain-label metamorphic, revision, and cross-environment testing.

Cycle 1 supports a narrow reproducibility statement: under the tested inputs and on the implemented core fields, a separate model reconstructed the permission-composition behavior from the frozen specification and produced outputs matching the creator-frozen expectations.

Cycle 1 did not validate the full output contract through the cross-model review. The reviewer intentionally omitted several specification-defined fields, including provenance and upper-bound outputs, after preregistering its interpretation of ambiguity in the specification. Cycle 1 also did not test through the cross-model review the full malformed-input boundary, revision-triggered recomputation, full provenance requirements, or the complete metamorphic suite.

3.4 Invariants directly exercised by the cross-model corpus

The cross-model corpus directly exercised behavior corresponding to I02 (non-expansion), I03 (resolved intersection), I09 (empty-set annihilation), I10 (non-indispensable bridges ignored), I11 (uncertainty preservation), I12 (determinate exclusion under unresolvedness), and I20 (no inferred hierarchy among arbitrary assertion symbols).

Other frozen invariants may be consistent with individual cases but were not subjected during Cycle 1 to their dedicated repeated-execution, permutation, regrouping, idempotence, restriction, malformed-input, provenance, revision, or metamorphic tests. For this reason, the paper distinguishes **invariant consistency** from **cross-model invariant testing**.

3.5 Cycle 1 closure and protocol limitation

An attempted subsequent blind counterexample phase did not produce admissible evidence because its preregistration artifacts were lost before reliable cryptographic sealing. The event was preserved as E12-P4-DISC-001 — Phase 4A artifact freeze failure. No Phase 4 counterexample was executed, the reference implementation was not revealed during the blind testing cycle, and the failed artifact process was not repaired retroactively.

E12 Review Cycle 1 was therefore closed with the disposition:

Survives blind cross-model reimplementations and held-out core testing.

This wording is a cross-model reproducibility disposition, not a claim of independent validation. Cycle 1 was model-to-model, tested a limited cross-model corpus, did not complete an admissible blind counterexample phase, and did not establish domain-specific validity, human reproducibility, or universal correctness.

4. Upstream Preconditions

4.1 Composition is conditional on its inputs

LiO v0.1 deliberately does not determine whether an assertion class is well defined, whether a bridge is genuinely indispensable, whether the declared derivation is admissible, or whether a bridge should permit a particular class of downstream assertion. Those judgments are supplied to the composition primitive. Deterministic composition therefore does not by itself guarantee scientifically legitimate authorization. The kernel can compose supplied permissions exactly while still producing an inappropriate downstream result if the representation, dependency declaration, or bridge adjudication supplied to it is defective.

The primitive answers

Given A , I_D , and $P(b)$, what authorization survives composition?

It does not answer what distinctions should appear in A , which bridges truly belong in I_D , what evidence should determine $P(b)$, or whether D is a valid derivation.

4.2 Assertion-space adequacy

Let A_D denote the finite assertion universe used for the declared derivation D (written A elsewhere for brevity). The first precondition is that A_D preserve every distinction whose loss could materially alter downstream authorization.

Suppose two scientifically different claims x and y are mapped to the same assertion symbol:

$$q(x) = q(y) = a.$$

Once that collapse occurs, LiO cannot recover the distinction. If a bridge permits a , the composition kernel cannot distinguish a weaker claim from a materially stronger one that has been encoded under the same symbol. Relevant distinctions may include association versus causation, possible versus established, population-specific versus universal, conditional versus unconditional, or contributory versus sufficient causal status.

This yields a representation constraint:

LiO cannot preserve a distinction that is absent from A_D .

The granularity of A_D therefore determines the semantic resolution at which LiO can constrain downstream claims. A coarse assertion universe can create assertion-class laundering: a stronger successor claim passes through because the representation has already erased the distinction that would have caused its exclusion.

The assertion universe need not be a globally fixed ontology for an entire scientific or technical domain. It may be constructed and versioned for a particular declared derivation. That scoping limits the blast radius of a representation failure but does not verify adequacy. If a materially relevant distinction is absent from A_D —including one recognized only after bridge adjudication has begun—LiO cannot represent or preserve that distinction within the current derivation without reopening or versioning its representation. Per-derivation assertion-space adequacy is therefore an upstream assertion supplied to, rather than verified by, LiO.

4.3 Semantic commensurability across bridges

Intersection is meaningful only when every bridge permission refers to a common assertion semantics. If two adapters use the same symbol while intending different meanings, then

$$P(b_1) \cap P(b_2)$$

may be mathematically well formed while being semantically incoherent.

All bridge permissions participating in a declared derivation must therefore already be mapped into a semantically commensurable assertion universe. Ontology-matching research demonstrates that semantic correspondence across heterogeneous representations is itself a substantive technical problem, and alignment repair is needed when mappings become incoherent [6,7]. LiO treats assertion symbols as opaque. It does not inspect their definitions or determine whether two domains use them equivalently. Cross-domain mapping is an upstream responsibility.

4.4 Local adjudication of bridge permissions

For a resolved bridge b , LiO receives a permission set $P(b) \subseteq A$. Membership must represent a positive upstream judgment that assertion class a is permitted to pass that inferential transition. The judgment must not be reverse-engineered from the desired final conclusion.

A future adapter can be represented abstractly as

$$G_D(b, a, E, C, V) \rightarrow \{\text{PERMIT}, \text{EXCLUDE}, \text{UNRESOLVED}\},$$

where E denotes evidence, C relevant context, and V the versioned adjudication rules. For a fully resolved bridge,

$$P(b) = \{a \in A : G_D(b, a, E, C, V) = \text{PERMIT}\}.$$

The adjudication must be local enough that the bridge-level permission is derived from evidence and inferential properties relevant to that bridge rather than from the desired downstream result. Otherwise the architecture becomes circular:

$$\text{desired conclusion} \rightarrow P(b) \rightarrow \text{same desired conclusion}.$$

LiO cannot detect this failure internally.

Because the adapter interface exposes the resulting bridge-local judgment for a given (b, a) pair, the composition layer has no operational signal that distinguishes legitimate bridge-local adjudication from backward rationalization that happens to return the same PERMIT value. This is an **adjudication circularity risk: a structural risk in the interface**. It does not make G_D logically circular by definition; it means non-

circularity must be established by the upstream adjudication procedure and its validation evidence rather than by the LiO kernel.

4.5 Reproducibility of adjudication

A deterministic kernel cannot compensate for an unstable upstream adjudication procedure. If two qualified adjudicators presented with the same evidence and rules assign materially different permission sets, downstream outputs will differ even though LiO itself remains deterministic.

End-to-end reproducibility therefore decomposes as

$$E \xrightarrow{G_D} P(b) \xrightarrow{LiO} Auth(a \mid D).$$

The first mapping is an empirical domain problem. Evidence-grading studies illustrate why this cannot be assumed: experienced reviewers can disagree substantially on complex evidence, while structured criteria and training can improve inter-rater reliability relative to unaided global judgment [14,15]. A future adapter should therefore expose decision criteria, provenance for each judgment, versioned rules, disagreements, and sufficient structure for independent replay. LiO v0.1 does not establish that such adjudication is reproducible; it makes the dependence explicit.

4.6 Context instantiation and local separability

The current primitive assumes that each resolved bridge can be represented by an individual permission set. That assumption fails when permission depends on context that has not yet been instantiated. For example, a rule of the form “assertion a is permitted only if condition c has independently been established” cannot be faithfully represented by a context-free $P(b)$ unless the relevant condition has first been resolved upstream.

Every resolved permission supplied to LiO must therefore already be fully instantiated for the context on which that permission depends.

A deeper boundary occurs when bridge permissions are irreducibly interactive. If the permission associated with b_1 cannot be determined independently of b_2 , the true object may be relational,

$$P(b_1, b_2),$$

rather than separable into $P(b_1)$ and $P(b_2)$. LiO v0.1 assumes sufficient local separability for bridge-wise permission sets to be meaningful. Where that assumption fails, the interaction must be represented upstream as a composite bridge or handled by a richer future formalism.

4.7 Completeness of the declared dependency set

LiO composes permissions across the bridges declared indispensable to a particular derivation. It does not discover whether that declaration is correct.

Let the true indispensable set be

$$I_D^i = \{b_1, b_2\},$$

with

$$P(b_1) = \{a\}, P(b_2) = \emptyset.$$

If the supplied derivation incorrectly declares only $I_D = \{b_1\}$, LiO will authorize a , even though the omitted bridge would have excluded it. Thus

LiO cannot detect a missing indispensable dependency.

The converse error also matters. If an irrelevant restrictive bridge is incorrectly declared indispensable, LiO may reject a claim that the actual derivation would permit. Under-declaration threatens authorization soundness relative to the intended derivation; over-declaration threatens completeness or usefulness.

4.8 Admissibility and alternative derivations

The current kernel does not establish that the supplied derivation is inferentially legitimate. A cyclic, unsupported, or otherwise inadmissible derivation can still contain internally consistent permission sets. LiO may compose them correctly while the larger argument remains defective. The declared derivation must therefore already be admissible under the inferential or domain-specific rules relevant to the application.

The primitive is also derivation-specific. It evaluates

$$Auth(a \mid D),$$

not global authorization across all possible derivations. If multiple alternative derivations exist, each would need to be represented separately under the current version. A higher-level rule combining alternative routes is outside LiO v0.1.

4.9 Consolidated preconditions

The current composition claim is conditional on six upstream requirements:

- C_1 : a declared derivation D ,
- C_2 : a complete and correctly supplied indispensable set I_D ,
- C_3 : legitimately adjudicated bridge permissions $P(b)$,
- C_4 : a derivation-scoped, semantically commensurable assertion universe A_D adequate to the material distinctions
- C_5 : permissions fully instantiated for relevant context and sufficiently separable,
- C_6 : an upstream-admissible derivation.

LiO does not verify any of C_1 – C_6 merely by receiving them as inputs. In particular, the adequacy of A_D , completeness of I_D , non-circularity and legitimacy of bridge adjudication, local separability, and admissibility of D remain upstream assertions subject to independent domain evaluation.

Given these conditions, the defensible core statement is:

Given a declared admissible derivation D , a complete supplied set of its indispensable inferential bridges, and upstream-adjudicated permissions over a derivation-scoped, semantically commensurable assertion universe, LiO v0.1 deterministically constrains a target assertion to the conjunction of the permission states supplied by every indispensable bridge, while preserving explicit unresolvedness.

This is narrower than a claim that LiO determines whether an assertion is scientifically true or globally justified. It states what the current primitive controls.

5. Limits and Non-Claims

5.1 Representation limits

Assertion-class laundering

Condition: Coarse or incomplete A_D collapses or omits materially distinct claims

Required response: Reopen/version A_D to preserve necessary distinctions

Status: Upstream: Representation

Semantic commensurability

Condition: Same symbol carries different meanings across bridges

Required response: Enforce common semantics before composition

Status: Upstream: Representation

LiO v0.1 does not define, validate, or guarantee the semantic adequacy of derivation-scoped A_D .

5.2 Adjudication limits**Adjudication circularity risk**

Condition: Legitimate bridge-local evaluation and backward-rationalized evaluation can expose the same $(b, a) \rightarrow \text{PERMIT}$ interface result

Required response: Validate bridge-local criteria, evidence, and decision procedure outside the composition kernel

Status: Upstream: Adjudication

Adjudication non-reproducibility

Condition: Same evidence produces materially different $P(b)$ without principled explanation

Required response: Standardize criteria; preserve disagreement; test reliability

Status: Upstream: Adjudication

Context non-instantiation

Condition: $P(b)$ depends on unresolved context

Required response: Resolve relevant context before supplying $P(b)$

Status: Upstream: Adjudication

LiO v0.1 does not adjudicate, validate, or guarantee the correctness or reproducibility of $P(b)$.

5.3 Structural and expressive limits**Missing indispensable dependency**

Condition: Supplied I_D omits a true indispensable bridge

Required response: Validate dependency completeness upstream

Status: Upstream: Derivation

Over-declaration

Condition: Supplied I_D includes a bridge not actually indispensable

Required response: Validate dependency minimality upstream

Status: Upstream: Derivation

Non-separable permissions

Condition: Joint permission cannot be decomposed into bridge-local sets

Required response: Represent a composite bridge or use richer formalism

Status: Future extension

Inadmissible derivation

Condition: Supplied derivation is unsupported or otherwise invalid under domain rules

Required response: Validate derivational admissibility upstream

Status: Upstream: Derivation

Joint assertion constraints

Condition: Individual assertion symbols are permitted but some combinations are not

Required response: Use relational or bundle-level representation if required

Status: Future extension

Conditional permissions

Condition: Permission depends on conditions not represented in fixed $P(b)$

Required response: Instantiate context upstream or enrich formalism

Status: Future extension

Alternative derivations

Condition: Authorization can proceed through multiple independent routes

Required response: Evaluate derivations separately or add a higher-level derivation combiner

Status: Outside v0.1

LiO v0.1 is not a general proof-graph evaluator and does not represent arbitrary Boolean dependency structures or every possible joint-permission policy.

5.4 Kernel and conformance limits

Output-contract ambiguity

Condition: Frozen specification defines fields omitted by the cross-model implementation

Required response: Resolve specification ambiguity in a future version without rewriting v0.1 evidence

Status: Specification issue

Schema/core conformance mismatch

Condition: Raw core acceptance can differ from schema rejection

Required response: Place a schema-validating conformance boundary before the mathematical kernel

Status: Conformance repair candidate

Downstream artifact retraction

Condition: Revision-triggered recomputation can change current authorization but v0.1 does not automatically retract or revalidate already-materialized downstream artifacts

Required response: Use an external dependency/revalidation mechanism where artifact retraction is required

Status: Outside v0.1

Cross-model canonicalization coverage

Condition: Cycle 1 did not exercise the complete canonicalization/representation-invariance contract through the cross-model review

Required response: Cover in a future externally controlled test cycle

Status: Not E12 Cycle 1 evidence

LiO v0.1 does not claim full conformance to every frozen output and input-boundary requirement on the basis of Cycle 1 alone.

5.5 Validation limits

Limited cross-model corpus

Condition: E12 Cycle 1 cross-model corpus contained 19 cases

Required response: Expand with an externally authored and controlled corpus, including real-domain cases

Status: Future validation

Model-to-model only

Condition: No independent human or laboratory replication

Required response: Add human and external implementation review

Status: Future validation

Subset of invariants directly exercised

Condition: Cycle 1 directly exercised I02, I03, I09–I12, and I20

Required response: Externally test remaining frozen invariants

Status: Future validation

Counterexample phase not completed

Condition: Phase 4 preregistration artifacts were lost before admissible execution

Required response: Run a future fresh cycle with external artifact custody

Status: Future validation

The nineteen E12 cases are not the full creator-side validation history. They are the cross-model evidence corpus for Cycle 1.

5.6 Consolidated non-claims

This paper does not claim novelty for set intersection, lattice meet, three-valued conjunction, permission propagation as a general idea, or dependency tracking as a general idea. It does not claim universal correctness, global authorization, correctness of domain-specific permission adjudication, adequacy of any particular semantic representation, full output-contract conformance, independent human validation, domain-specific scientific validity, or handling of alternative derivations within the current kernel.

The strongest empirical statement authorized by Cycle 1 is narrower: a blind cross-model reimplementations matched all seven public cases and all twelve creator-frozen held-out cases on the two implemented core fields, with no semantic disagreements on that corpus.

6. Residual Research Questions

6.1 What remains after composition is reduced

The formal reduction of LiO v0.1 makes the remaining research problem easier to locate. Once bridge permissions have been supplied over a common assertion universe, the composition mechanism is small:

$$\text{Auth}(a \mid D) = \bigwedge_{b \in I_D} v_b(a).$$

The difficult questions therefore occur largely before this conjunction is evaluated. The research program separates into

Representation → Adjudication → Composition.

The present executable primitive addresses the third component. The first two remain open.

6.2 RQ-001: What semantic distinctions must be represented?

LiO can protect only distinctions that are represented. If two materially different claims are encoded as the same assertion class, no subsequent composition rule can restore the distinction. The first research question is therefore:

Which meaning-bearing distinctions must an assertion representation preserve so that downstream permission composition does not authorize a materially stronger successor claim?

Candidate distinctions include epistemic strength, scope, modality, causality, temporal qualification, population, conditionality, authority, and evidential status. The present paper does not establish that this list is complete or that these dimensions should be represented in any particular formalism.

A future representation could take the form of typed semantic state,

$$S(x) = \{d_i : v_i\},$$

with transformations represented as

$$T : S(x) \rightarrow S(x')$$

and semantic change as

$$\Delta_T = S(x') \ominus S(x).$$

This is a research hypothesis, not part of the validated LiO v0.1 primitive. The empirical question is whether consequential transformations such as narrowing, expansion, strengthening, qualification loss, reversal, or ambiguity collapse can be represented with sufficient reliability for independent use.

6.3 RQ-002: How should semantic or inferential change produce permission?

Even a sufficiently expressive assertion universe does not determine what a bridge should permit. The second open problem is the mapping

$$G_D(b, a, E, C, V) \rightarrow \{\text{PERMIT}, \text{EXCLUDE}, \text{UNRESOLVED}\}.$$

The research question is:

Can domain-specific procedures reproducibly map the evidentiary and semantic state of an indispensable inferential transition to downstream assertion permissions?

A successful adapter would need more than a plausible mapping. Research on evidence claims and warrants already shows that claim justification depends on explicit standards linking evidence to conclusions [10,11], and contemporary work explicitly frames evidence as licensing some scientific claims while withholding others [13]. A future LiO adapter would therefore need explicit decision criteria, versioned rules, provenance for each judgment, independent replay, and preserved disagreement. The distinction is fundamental:

$$\text{deterministic composition} \Rightarrow \text{reproducible adjudication}.$$

LiO v0.1 establishes the former under tested conditions. A complete scientific system would require evidence for both.

6.4 RQ-003: Where does the powerset permission model stop being adequate?

The current primitive associates each resolved bridge with $P(b) \subseteq A$. This representation assumes that permission can be expressed as the individual admissibility of assertion symbols. It cannot directly represent every joint or conditional policy.

For example, a bridge may permit a and b individually but prohibit asserting them jointly. A simple permission set $P(b) = \{a, b\}$ cannot encode that restriction. Likewise, a permission may be conditional on another established fact, or the permission of one bridge may depend irreducibly on the state of another.

The research question is:

Under what classes of inference-control problem is bridge-wise powerset permission sufficient, and when are conditional, relational, or higher-order permission structures required?

These cases define the expressive boundary of the current unary, context-instantiated, locally separable kernel. They do not falsify that kernel within its declared scope.

6.5 The deeper coupling

The three research problems are coupled. A future semantic representation may detect that a transformation changes a claim from “X may contribute to Y under condition Z” to “X causes Y.” That transformation could be represented as

$$S(x) \rightarrow \Delta_T \rightarrow S(x').$$

Adjudication would then ask what that measured transformation permits,

$$\Delta_T \rightarrow P(b),$$

and LiO composition would operate afterward,

$$P(b_1), P(b_2), \dots, P(b_n) \rightarrow \bigcap_{b \in I_D} P(b).$$

This produces the larger candidate architecture

$$S(x) \rightarrow \Delta_T \rightarrow P(b) \rightarrow \bigcap_{b \in I_D} P(b).$$

Only the final operation has reached the frozen executable validation state reported in this paper. The upstream stages remain research hypotheses.

6.6 Language-mediated transformations

The representation problem becomes especially visible when information is repeatedly transformed through language. A downstream statement can remain lexically similar to its source while changing scope, evidential strength, causal status, qualification, authority, or conditionality. A system may therefore preserve substantial textual content while failing to preserve relations that mattered to the original claim.

This motivates a broader hypothesis: semantic continuity may depend less on preserving words than on preserving the meaning-bearing relations required by the downstream use of those words. LiO v0.1 does not establish that hypothesis. It provides a possible terminal constraint mechanism if such transformations can eventually be represented and adjudicated.

6.7 Research sequence

The next research stages should proceed conservatively. First, determine whether consequential semantic distinctions can be represented reproducibly. Second, determine whether measured transformations can produce reproducible bridge permissions. Third, test whether the current composition primitive remains sufficient when those permissions come from real scientific or technical domains rather than synthetic adapters. Only demonstrated failures should justify enlarging the kernel.

The residual question is therefore more precise than a broad claim about semantic understanding:

Can meaning-bearing transformation be measured well enough to determine what a downstream inference is still allowed to say?

7. Discussion and Conclusion

LiO v0.1 is deliberately smaller than the architecture that motivated it. The current mechanism does not discover hidden relations, define semantic operators, adjudicate scientific evidence, validate cross-domain equivalence, or decide truth. It composes already-adjudicated permissions across a supplied set of indispensable inferential bridges and preserves unresolvedness when authorization cannot yet be determined.

That reduction is not a retreat from the broader research problem. It identifies where the present evidence is strongest and where the remaining scientific work actually lies. The mathematical kernel is familiar: set intersection and coordinatewise three-valued conjunction. The unresolved questions concern representation, adjudication, dependency completeness, semantic commensurability, and expressive scope. Separating those problems prevents the simplicity of the kernel from being confused with the difficulty of the larger system.

E12 Review Cycle 1 provides a limited but useful replication result. A separate model reconstructed the frozen mechanism from specification, reproduced seven public cases, and then matched twelve creator-frozen held-out cases exactly on the two implemented core fields. The review also exposed specification and artifact-custody weaknesses that were preserved rather than erased. The resulting evidence supports a narrow claim of cross-implementation reproducibility on the tested corpus, not universal correctness or independent scientific validation.

The current contribution can therefore be stated without relying on mathematical novelty:

Given a declared admissible derivation, a complete supplied set of indispensable inferential bridges, and upstream-adjudicated permissions over a derivation-scoped, semantically commensurable assertion universe, LiO v0.1 deterministically constrains downstream assertion authorization to what every indispensable bridge permits, while preserving explicit unresolvedness.

The next scientific question lies upstream of that result. If meaning-bearing transformations can be represented and bridge permissions can be adjudicated reproducibly, LiO offers a small composition layer that prevents downstream authorization from silently exceeding those judgments. If those upstream problems cannot be solved reliably, the kernel does not rescue them. That boundary is part of the result.

Appendix A. Claim Ledger

A.1 Established by E12 Review Cycle 1

C-CORE-001 — Cross-model reproducibility on the Cycle 1 corpus. The frozen LiO v0.1 permission-composition behavior produced matching outputs across the creator implementation and the separate-model implementation on the tested public and held-out cases for the implemented core fields. Evidence: 7/7 public exact and 12/12 held-out exact on `exact_final_permissions` and `target_authorized`.

C-CORE-002 — Resolved and partially unresolved composition matched creator-frozen expectations on the held-out corpus. No semantic disagreement was observed across H001–H012 on the implemented fields.

C-CORE-003 — Explicit unresolvedness was preserved in the cross-model implemented behavior on cases exercising that state. Cycle 1 supports this only within the tested corpus.

A.2 Conditional upstream dependencies

C-UPSTREAM-001 — Assertion-space adequacy. The assertion universe must preserve semantic distinctions whose loss would materially change authorization. This is required but not validated.

C-UPSTREAM-002 — Reproducible, non-circular bridge adjudication. Permission sets must be generated by explicit upstream procedures rather than reverse-engineered from desired conclusions. This is required but not validated.

C-UPSTREAM-003 — Local separability and context instantiation. Bridge-local permission sets must be meaningful for the declared context. This is required but not validated.

C-UPSTREAM-004 — Dependency completeness and derivational admissibility. The supplied indispensable bridge set and declared derivation must be correct under upstream rules. This is required but not validated by the kernel.

A.3 Explicit non-claims

LiO v0.1 does not claim novelty of set intersection, lattice meet, three-valued conjunction, generic permission propagation, or generic dependency tracking. It does not claim full output-contract conformance from E12 Cycle 1, correctness of domain-specific permission adjudication, adequacy of any particular semantic representation, universal correctness, global authorization, alternative-derivation handling, independent human or laboratory validation, or domain-specific scientific validity.

A.4 Open research questions

RQ-001: Which semantic distinctions must be represented so that downstream composition does not launder materially stronger claims through coarse encoding?

RQ-002: Can evidence and semantic transformation be mapped reproducibly to bridge-local permission states?

RQ-003: When is bridge-wise powerset permission sufficient, and when are relational, conditional, or higher-order permission structures required?

References

-
- [1] Doyle, J. (1979). *A Truth Maintenance System*. **Artificial Intelligence**, 12(3), 231–272. [https://doi.org/10.1016/0004-3702\(79\)90008-0](https://doi.org/10.1016/0004-3702(79)90008-0)
- [2] Martins, J. P., & Shapiro, S. C. (1988). *A Model for Belief Revision*. **Artificial Intelligence**, 35(1), 25–79. [https://doi.org/10.1016/0004-3702\(88\)90031-8](https://doi.org/10.1016/0004-3702(88)90031-8)
- [3] Cousot, P., & Cousot, R. (1979). *Systematic Design of Program Analysis Frameworks*. In **Proceedings of the 6th ACM SIGACT-SIGPLAN Symposium on Principles of Programming Languages**, 269–282. <https://doi.org/10.1145/567752.567778>
- [4] Fitting, M. (1991). *Kleene’s Logic, Generalized*. **Journal of Logic and Computation**, 1(6), 797–810. <https://doi.org/10.1093/logcom/1.6.797>
- [5] Buneman, P., Khanna, S., & Tan, W.-C. (2001). *Why and Where: A Characterization of Data Provenance*. In **Database Theory — ICDT 2001**, Lecture Notes in Computer Science 1973, 316–330. https://doi.org/10.1007/3-540-44503-X_20
- [6] Shvaiko, P., & Euzenat, J. (2013). *Ontology Matching: State of the Art and Future Challenges*. **IEEE Transactions on Knowledge and Data Engineering**, 25(1), 158–176. <https://doi.org/10.1109/TKDE.2011.253>
- [7] Santos, E., Faria, D., Pesquita, C., & Couto, F. M. (2015). *Ontology Alignment Repair through Modularization and Confidence-Based Heuristics*. **PLOS ONE**, 10(12), e0144807. <https://doi.org/10.1371/journal.pone.0144807>
- [8] Arden, O., Liu, J., & Myers, A. C. (2015). *Flow-Limited Authorization*. In **2015 IEEE 28th Computer Security Foundations Symposium**, 569–583. <https://doi.org/10.1109/CSF.2015.42>
- [9] Melicher, W., Xu, A., Zhao, V., Potanin, A., & Aldrich, J. (2022). *Bounded Abstract Effects*. **ACM Transactions on Programming Languages and Systems**, 44(1), 1–48. <https://doi.org/10.1145/3492427>
- [10] Gough, D. (2021). *Appraising Evidence Claims*. **Review of Research in Education**, 45(1), 1–26. <https://doi.org/10.3102/0091732X20985072>
- [11] Ketokivi, M., & Mantere, S. (2021). *What Warrants Our Claims? A Methodological Evaluation of Argument Structure*. **Journal of Operations Management**, 67(6), 755–776. <https://doi.org/10.1002/joom.1137>
- [12] Herd, B., Kelly, J., Sabsch, J., & Gauerhof, L. (2026). *Towards a Compositional Semantics for Quantitative Confidence Assessment in Assurance Arguments*. **arXiv preprint arXiv:2605.22213**. <https://doi.org/10.48550/arXiv.2605.22213>
- [13] Li, H. (2026). *The Calibration Turn in AI-Assisted Research: A Conceptual and Methodological Framework for Evidence-Licensed Claims*. **arXiv preprint arXiv:2606.31273**. <https://doi.org/10.48550/arXiv.2606.31273>
- [14] Berkman, N. D., Lohr, K. N., Morgan, L. C., Kuo, T.-M., & Morton, S. C. (2013). *Interrater Reliability of Grading Strength of Evidence Varies with the Complexity of the Evidence in Systematic Reviews*. **Journal of Clinical Epidemiology**, 66(10), 1105–1117.e1. <https://doi.org/10.1016/j.jclinepi.2013.06.002>
- [15] Mustafa, R. A., Santesso, N., Brožek, J., et al. (2013). *The GRADE Approach Is Reproducible in Assessing the Quality of Evidence of Quantitative Evidence Syntheses*. **Journal of Clinical Epidemiology**, 66(7), 736–742.e5. <https://doi.org/10.1016/j.jclinepi.2013.02.004>
- [16] OASIS. (2013). *eXtensible Access Control Markup Language (XACML) Version 3.0*. **OASIS Standard**, 22 January 2013. <https://docs.oasis-open.org/xacml/3.0/xacml-3.0-core-spec-os-en.html>
- [17] Ramli, C. D. P. K., Nielson, H. R., & Nielson, F. (2014). *The Logic of XACML*. **Science of Computer Programming**, 83, 80–105. <https://doi.org/10.1016/j.scico.2013.05.003>
- [18] W3C. (2007). *Web Services Policy 1.5 — Framework*. **W3C Recommendation**, 4 September 2007. <https://www.w3.org/TR/2007/REC-ws-policy-20070904/>
-

Evidence-control note

This citation pass was conducted separately from mechanism development and E12 validation. Bibliographic metadata and source relevance were checked against live academic-research indexes before insertion. References [12] and [13] are explicitly identified as arXiv preprints rather than peer-reviewed journal articles; [16] and [18] are normative technical standards rather than peer-reviewed research articles. The related-work language is subtractive: cited overlap is used to narrow claims, not to infer novelty from gaps in the retrieved literature. This bounded pass does not establish exhaustive prior-art coverage or absence of an exact predecessor.