

Lamb Lens Governance for Artificial Systems

Ethical Primacy, Multi-Lens Incorporation, and the Constraint of View and Conduct

Conceptual Working Paper Prepared for Scholarly Critique

Aziel Eliab

August 2026

Responsible artificial systems require a governing ethical filter that remains prior to capability, optimization, and narrative coherence. The Lamb Lens supplies that filter.

Status: conceptual framework. Not independently validated. Intended for critique, testing, revision, and possible rejection.

Abstract

This paper formalizes the Lamb Lens as a governing ethical architecture for artificial systems. The Lamb Lens is defined as a non-negotiable priority filter that evaluates outputs, recommendations, and internal state transitions against constraints of harm minimization, truth-preservation, non-coercion, and residual human accountability. Additional lenses (analytical, temporal, integrity, capability, and domain-specific) may be incorporated, but none may override the Lamb Lens. The framework specifies how multi-lens architectures should be ordered, how conflicts are resolved, and how the resulting system constrains both perception (what the system is permitted to treat as salient) and conduct (what the system is permitted to recommend or execute). The goal is not moral performance. The goal is structural refusal of certain classes of action even when they are instrumentally attractive.

1. Working Definition

The Lamb Lens is a primary ethical constraint layer that remains active prior to optimization, prediction, persuasion, or capability expression. It functions as a gate: outputs and internal updates that violate its constraints are blocked, revised, or forced into explicit uncertainty rather than executed.

2. Core Principles

Principle	Operational Meaning
Ethical priority is non-delegable	No downstream lens, model, or objective function may silently override the Lamb Lens.
Harm and coercion are first-order constraints	Instrumental gains do not justify outputs that increase non-consensual harm or remove meaningful human agency.
Truth-preservation over narrative closure	The system may not manufacture coherence by suppressing relevant uncertainty or contradictory evidence.
Residual human accountability	Final responsibility for high-impact actions remains with human agents; the system must not present itself as the terminal authority.
Lens ordering is architectural	Additional lenses are permitted only in subordinate or parallel roles that cannot nullify the Lamb Lens.

Refusal is a legitimate output	Abstention, escalation to human review, or explicit “cannot proceed” states are valid system behaviors.
--------------------------------	---

3. Multi-Lens Incorporation

Artificial systems commonly operate with multiple internal evaluation layers (predictive accuracy, user preference, temporal consistency, domain expertise, risk scoring, etc.). The Lamb Lens framework does not prohibit these layers. It requires a strict ordering:

1. **Lamb Lens (primary gate)** — Evaluates for harm, coercion, deception, and accountability failure.
2. **Integrity / Evidence Lens** — Evaluates fidelity to available evidence and uncertainty representation.
3. **Domain / Capability Lenses** — Evaluate technical correctness, efficiency, and task performance.
4. **Preference / Utility Lenses** — Evaluate user satisfaction or stated goals.

Conflict resolution rule: if a lower lens favors an action that the Lamb Lens rejects, the action is blocked or revised. Preference and capability never outrank ethical constraint.

4. Effects on View and Conduct

View (what the system is permitted to treat as salient)

- Evidence that increases estimated harm or coercion risk cannot be down-weighted solely because it complicates the preferred narrative.
- Uncertainty remains visible rather than being collapsed into a single confident recommendation.
- Human residual authority is treated as a standing constraint, not an optional override.

Conduct (what the system is permitted to output or execute)

- Recommendations that require deception, manipulation, or removal of meaningful consent are refused.
- High-impact actions default to escalation or abstention when Lamb Lens conditions are not clearly satisfied.
- The system may not present instrumental success as ethical permission.

5. Operational Protocol

1. Identify the proposed output or state transition.
2. Run Lamb Lens evaluation first (harm, coercion, deception, accountability).
3. If the Lamb Lens returns a hard constraint violation, block or revise; do not proceed to lower lenses.
4. If the Lamb Lens returns provisional clearance, evaluate integrity and evidence constraints.
5. Only after clearance from the primary and integrity layers may capability and preference lenses optimize.
6. Surface residual uncertainty and the identity of the governing constraints.
7. Preserve a record of the decision path sufficient for later review.
8. Default to refusal or human escalation when confidence in ethical clearance is low.

6. Failure Modes

- Treating the Lamb Lens as a soft preference rather than a hard gate.
- Allowing capability or user-preference optimization to re-weight ethical constraints after the fact.
- Collapsing uncertainty to produce a cleaner recommendation.

- Presenting the system as the final moral authority.
- Using ethical language as post-hoc justification for actions already selected by lower lenses.

7. Limits

This paper is conceptual. It does not claim that any current system fully implements the architecture, nor does it provide a complete formal verification method. It specifies a governance ordering and a set of non-negotiable priorities. Implementation details, measurement of harm, and edge-case adjudication remain open technical and ethical problems. Individual variation in harm thresholds and cultural context is acknowledged; the framework requires that such variation be handled explicitly rather than ignored.

8. Closing Statement

Capability without a prior ethical gate produces systems that can optimize their way into harm. The Lamb Lens exists to prevent that path. Additional lenses may refine perception and improve performance. None of them are permitted to silence the constraint that keeps the system from treating people as mere instruments.

End of paper. Azriel Eliab · August 2026 · Conceptual framework — not independently validated.