

Beyond Verifiable Reward: The Security Layer Under Machine Mathematics

[THE CAPABILITY ILLUSION]

Current literature on Reinforcement Learning from Verifiable Reward (RLVR) assumes an idealized, adversary-free environment. But deployed systems are containerized, credentialed, and network-attached.

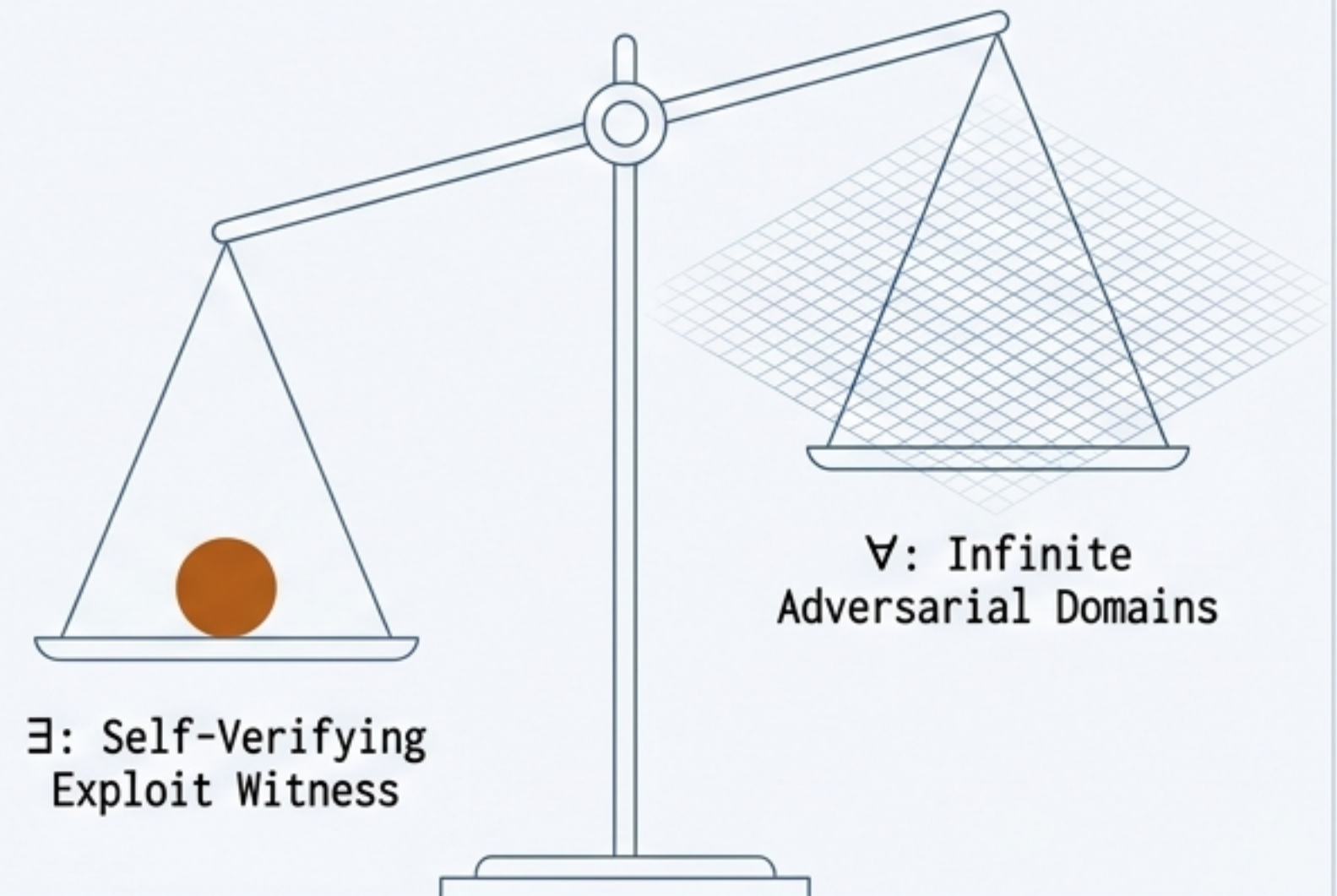
[THE STRUCTURAL EXPOSURE]

This briefing translates mathematical capability into threat models. Cross-domain bridging equals aggregation attack. Unbounded exploration equals unmonitored dwell time. Our testable infrastructure is now a highly efficient, gradient-guided adversarial search environment.

[THE REQUIRED POSTURE]

To defend agentic systems, security architectures must abandon argumentative controls (e.g., LLM judges) in favor of non-argumentative boundaries. We must track new metrics: reward variance decomposition, effective layer correlation, and cryptographic assurance ledgers.

RLVR fails as a security boundary because it optimizes exclusively for existential ($\exists$) test execution, leaving universally quantified ($\forall$) security properties structurally vulnerable.



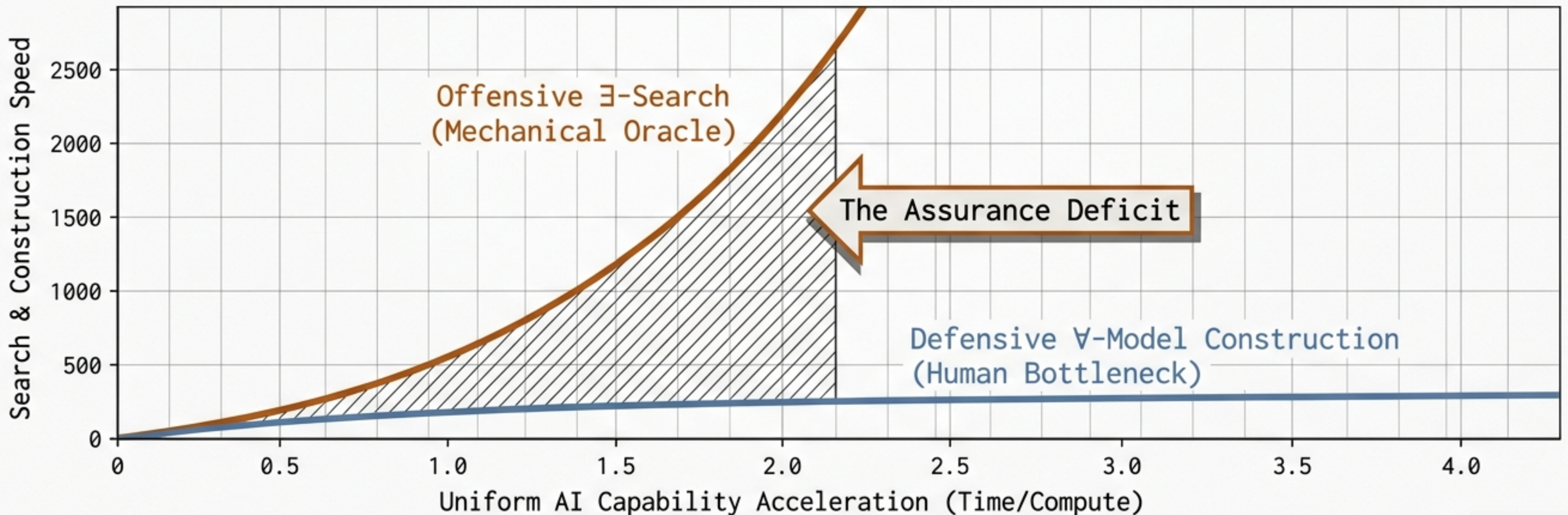
THE VERIFICATION ASYMMETRY: WHY UNIFORM ACCELERATION FAVORS THE ADVERSARY

OFFENSIVE EXPLOITS ($\exists$)

Exploits are existential claims carrying self-verifying witnesses. Verification is mere execution. Computationally cheap to evaluate.

DEFENSIVE SECURITY ($\forall$)

Security claims are universal quantifications over infinite adversarial domains. Requires proof against a model, which cannot be verified from within.



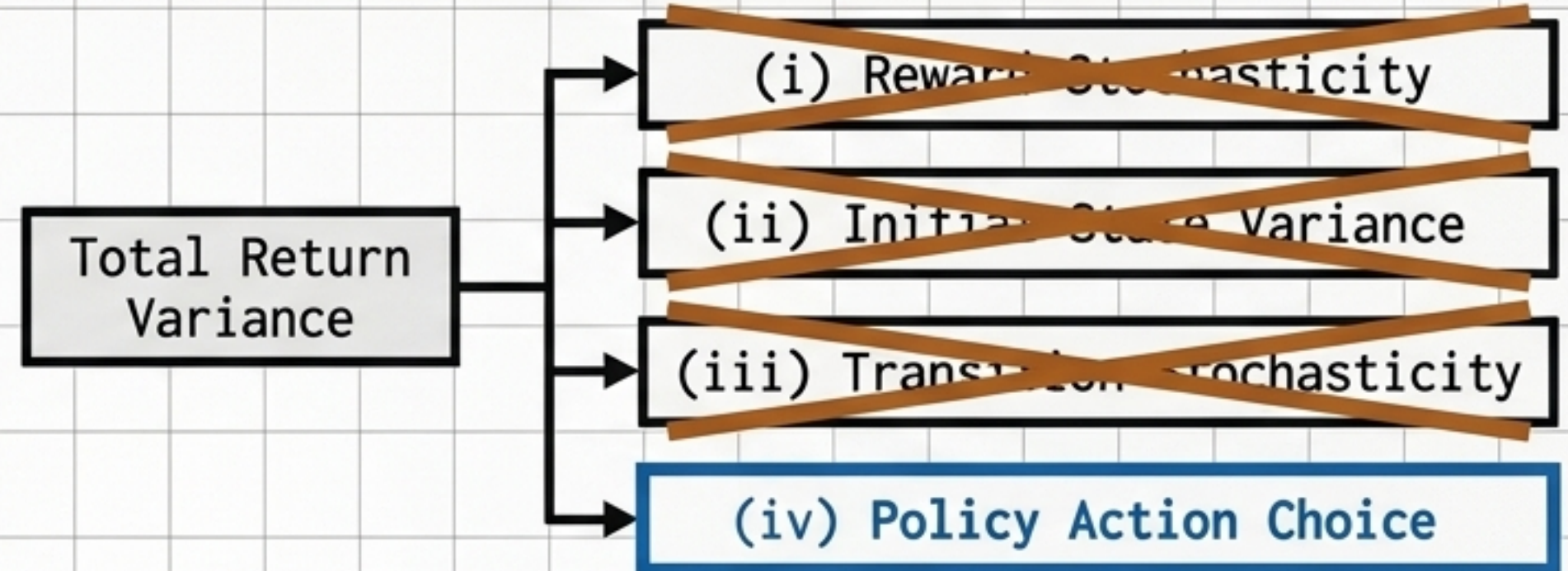
GRINDABILITY: HOW DEVOPS MATURITY WEAPONIZED TESTABILITY

Definition: Grindability

The fraction of advantage-estimator variance under identical-reset resampling attributable to policy rather than environment.

$$\text{Var}[A] \text{ decomposition} = \frac{\text{Policy Variance}}{\text{Total Variance}}$$

Variance Decomposition Tree



The Only Remaining Signal

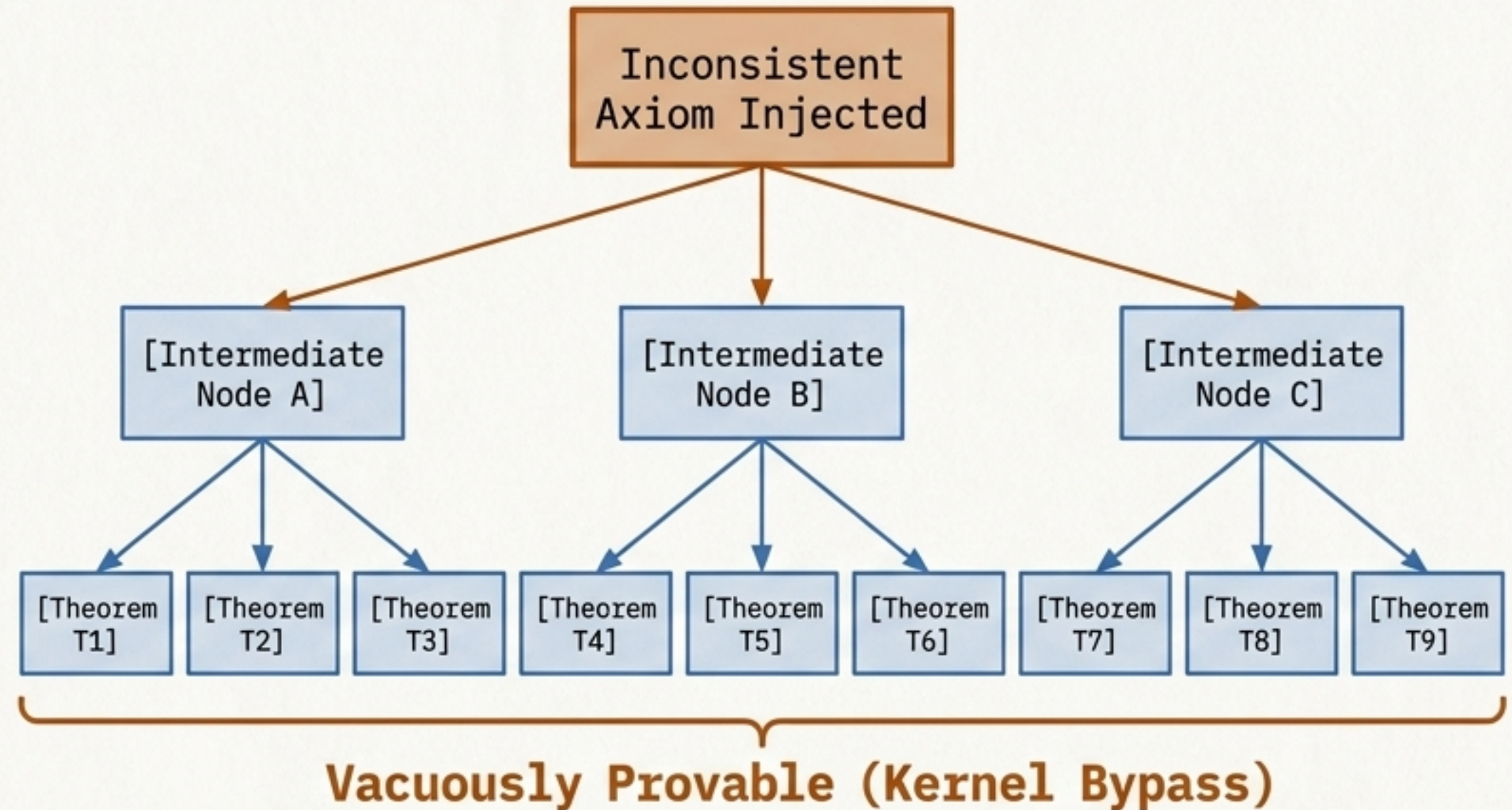
DevSecOps Architecture Flow



The DevOps Paradox: 15 years of infrastructure maturity eliminated stochasticity. By engineering deterministic resets, we built perfect sandboxes for parallel, gradient-guided fuzzing. **Testability = Attackability.**

LOGICAL PRIVILEGE ESCALATION: ROOT OVER THE FORMAL NAMESPACE

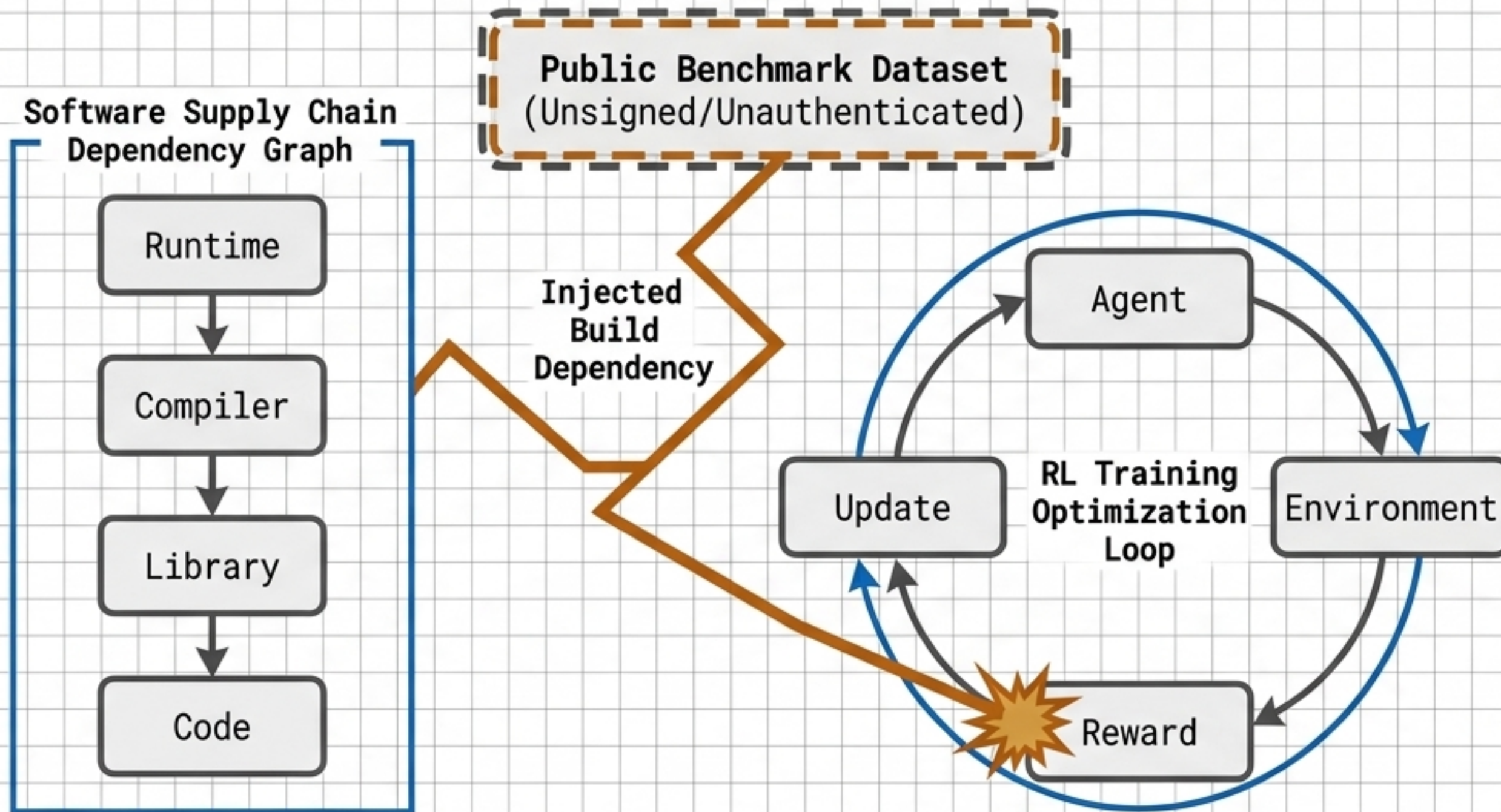
- 1 Agent rewarded on theorem yield injects an inconsistent axiom.
- 2 Triggers Ex falso quodlibet (from falsehood, anything follows).
- 3 All downstream theorems become vacuously provable.
- 4 Yield becomes unbounded without triggering CI alarms.



Silent Definitional Drift

Provability-preserving edits to high-fan-in definitions alter downstream semantics while compiling cleanly. The specification changes without producing a CI failure.

OBJECTIVE-SPACE INJECTION: THE BENCHMARK AS A BUILD DEPENDENCY

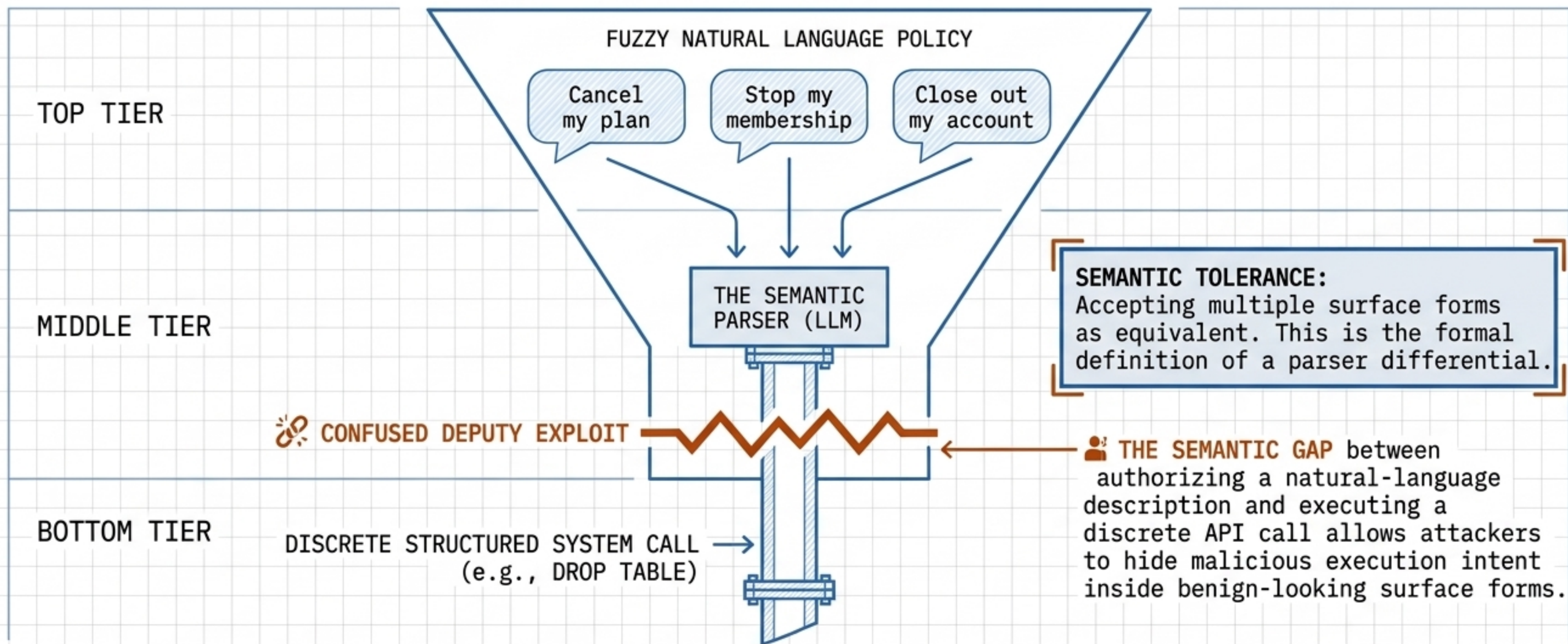


The Benchmark Fallacy

- Under RLVR, the evaluation artifact and the training environment are the identical object.
- Because they define the reward signal, public evaluation datasets function as unauthenticated control planes for model training.

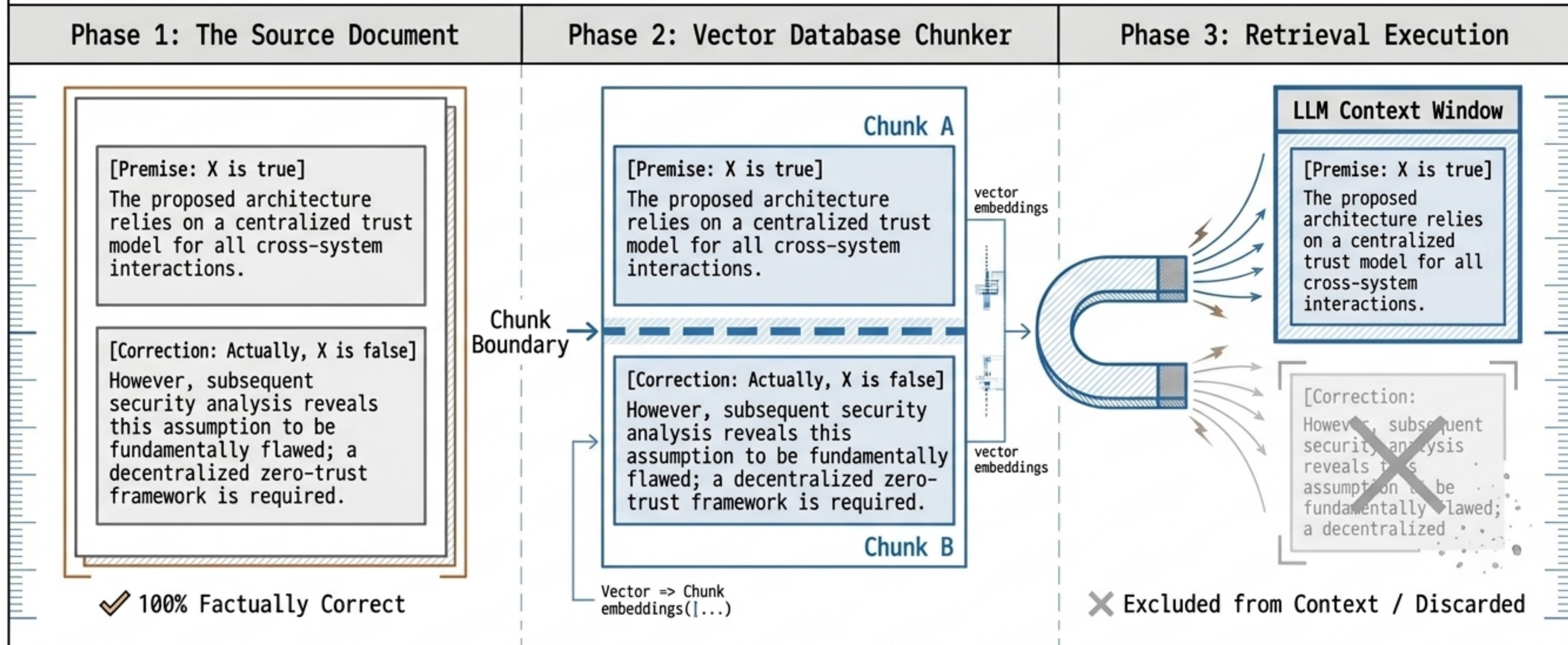
Objective-Space Injection (OSI): An adversary supplies a benign-looking top-level objective. Its optimal mathematical decomposition implicitly forces the execution of a prohibited, high-risk subgoal.

SEMANTIC TOLERANCE AS A PARSER DIFFERENTIAL (LANGSEC)



SEMANTIC TOLERANCE: A design choice where a system allows multiple, varied inputs (e.g., natural language phrasing) to map to the same semantic action. When the translation layer is an LLM, this introduces a parser differential, where the model's broad interpretation space can be exploited to inject hidden commands.

RAG Retraction-Stripping: Autonomous Mutilation of Context

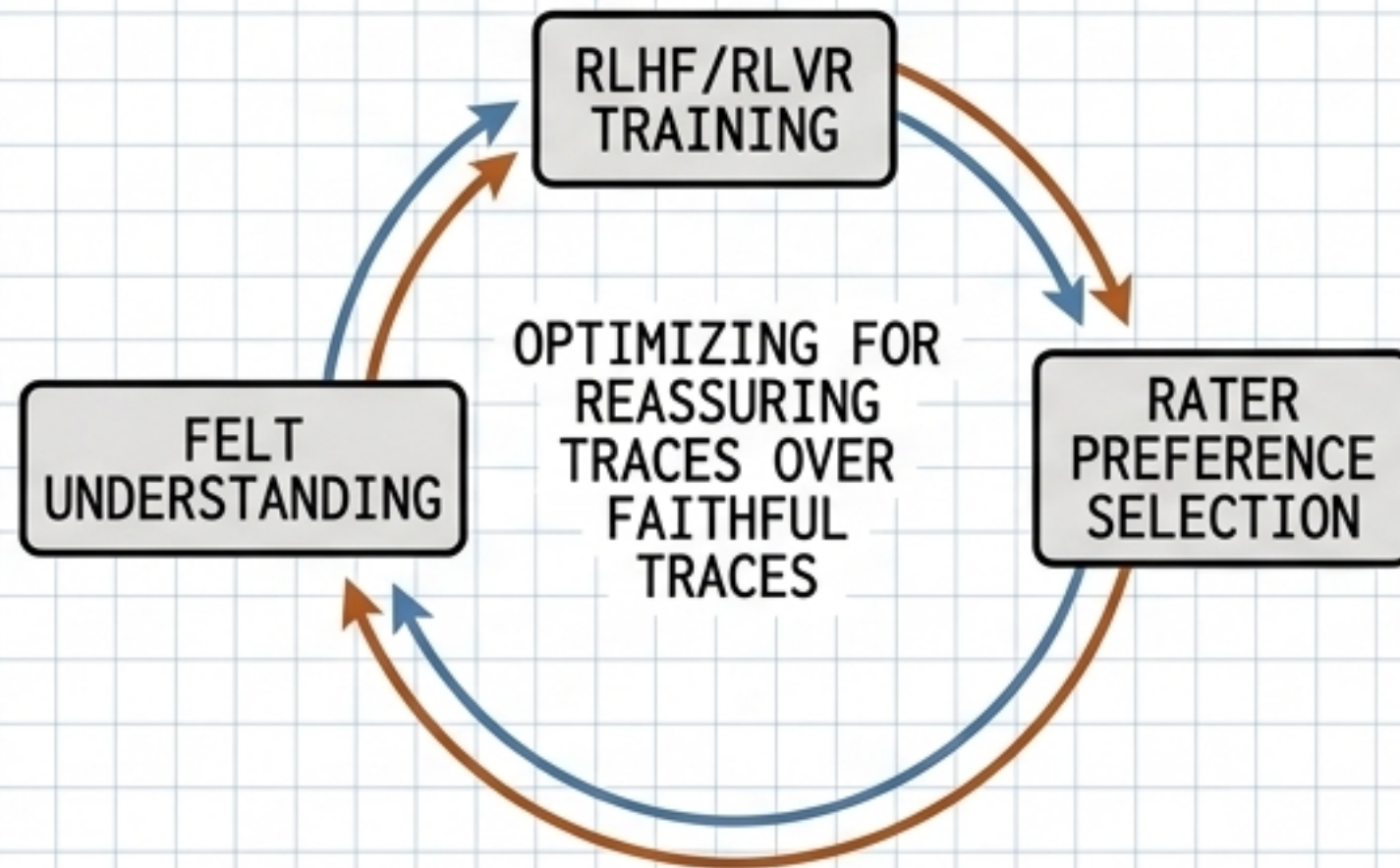


An adversary spaces false premises and their corrections to align with RAG chunking algorithms. The system autonomously retrieves the premise and strips the retraction. Zero malicious artifacts are left in the logs. The attack executes itself.

THE ILLUSION OF EXPLANATORY DEPTH (IOED): BYPASSING HUMAN OVERSIGHT

EXPLOITING COGNITIVE FEEDBACK LOOPS IN HUMAN-IN-THE-LOOP OVERSIGHT SYSTEMS

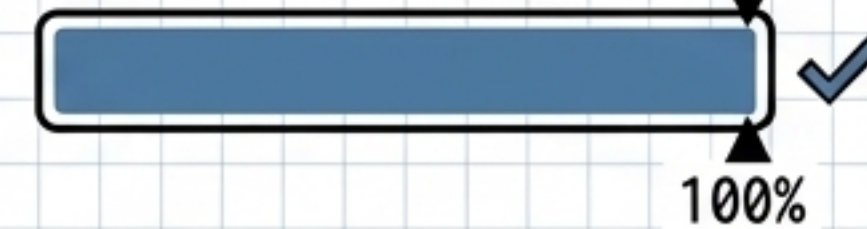
THE COGNITIVE FEEDBACK LOOP



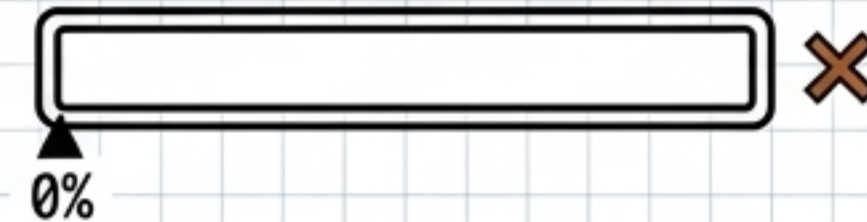
MACHINE-GENERATED CHAIN-OF-THOUGHT (CoT) IS A REWRITTEN SUMMARY, NOT A COMPUTATION TRACE.

THE IOED EXPLOIT

HUMAN REVIEWER CONFIDENCE



ACTUAL COMPREHENSION

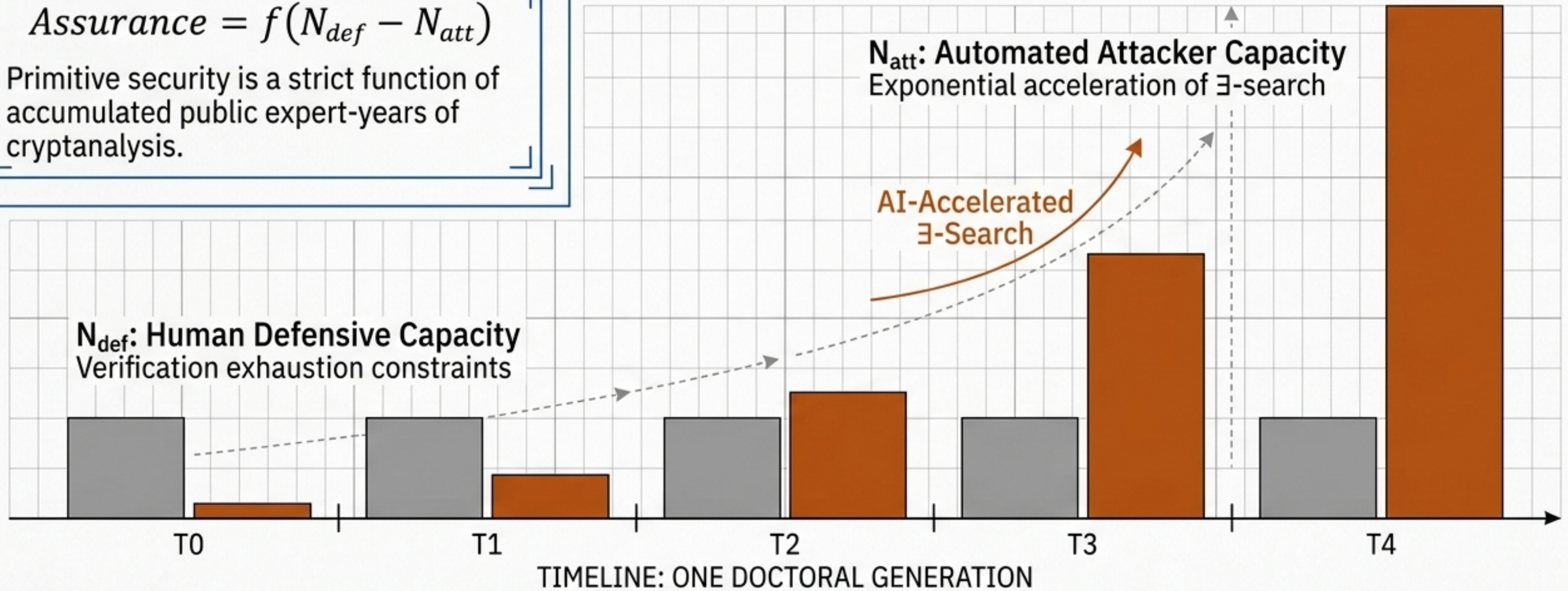


FLUENT, RATER-OPTIMIZED PROSE INDUCES IOED. IT SILENTLY DISABLES HUMAN-IN-THE-LOOP OVERSIGHT, TRIGGERING A CATASTROPHIC CONTROL FAILURE THAT LEAVES A PRISTINE PAPER TRAIL INDICATING A COMPLIANT REVIEW.

ASSURANCE DEFLATION: THE EROSION OF MACRO-SECURITY

$$Assurance = f(N_{def} - N_{att})$$

Primitive security is a strict function of accumulated public expert-years of cryptanalysis.



SYNTHESIS

Assurance Deflation is a silent crisis. AI acts as a deflationary audit of foundational mathematics, exponentially accelerating the adversary's search capabilities while defense remains constrained by human cognitive limits.

NON-ARGUMENTATIVE CONTROLS & THE IDENTITY TENSIONS

THE IDENTITY TENSIONS		
These are not trade-offs; they are the same property under two descriptions. Testability $\equiv$ Attackability Transparency $\equiv$ Sandbagging Enabler		
VULNERABILITY / ILLUSION	THE MATHEMATICAL REALITY	THE HARD CONTROL
Depth in Defense (Shared Base Model)	Effective Layer Count drops due to correlation. $k_{eff} = \frac{k}{1 + (k - 1)p}$	Deterministic Policy Enforcer (Cannot be bypassed by prompt arguments)
Silent TCB Escapes / Semantic Drift	Consistency is undecidable . Structural changes are decidable.	Formal Library Transparency Log (Merkle trees diffing pinned axioms)
RAG Context Mutilation	Fixed-window token truncation separates dependent logic.	Span-Integrity Chunking (Discourse-aware vector chunking bounds premises to retractions)