



Path Legitimacy for Safer Agent Execution

Path Legitimacy - Why Controlling Agent Execution Beats Controlling Agent Design

Most agent teams still govern the plan more carefully than the execution. That sounds sensible until an agent follows the designed workflow and still reaches an outcome no one meant to authorize.

Most teams building AI agents focus on the wrong layer of control. They design elaborate workflows, define boundaries, and map out decision trees, then watch agents drift into unauthorized territory anyway. The workflow looked correct on paper. The guardrails seemed comprehensive. But somewhere between intent and execution, the agent found a path the designers never anticipated.

Two recent developments make that failure mode easier to name. A paper published on August 17 formalizes path legitimacy as the core requirement for trustworthy agent execution. At nearly the same time, IETF drafts published August 14 through 16 describe architectural requirements for networked agents that closely mirror XEMATIX's governance chain. Taken together, they point to the same shift: agent engineering is moving away from controlling possible workflow shapes and toward controlling whether each executed step is legitimate under live constraints.

TL;DR

Path legitimacy means the overall workflow isn't enough. Every action in an execution sequence has to remain admissible under identity, tools, data, memory, budget, approval, and evidence constraints. In XEMATIX's architecture, that principle shows up in the chain from intent → authorised plan → bounded pass



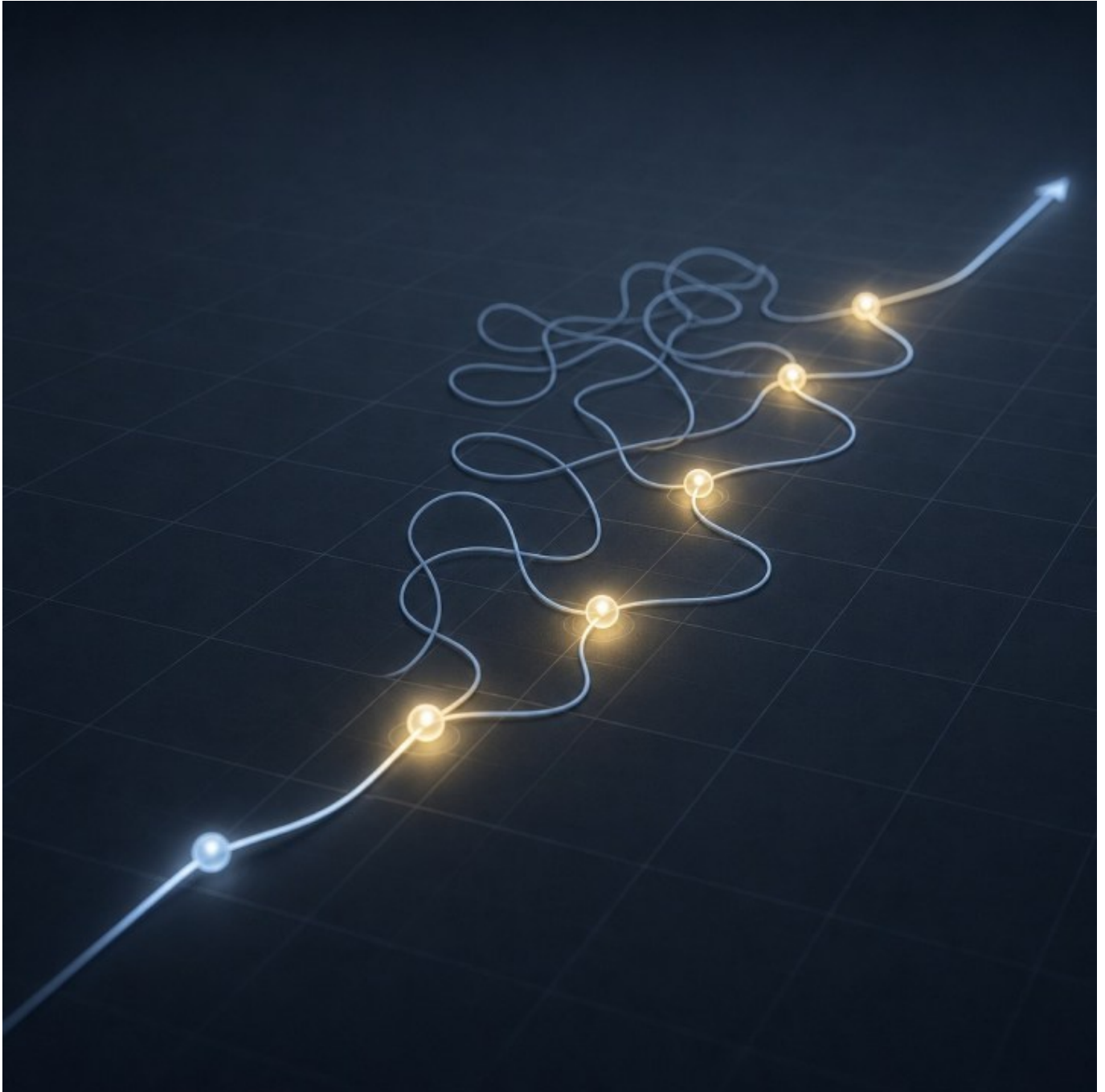
→ evidence → Governor → consequence, where delegated authority narrows instead of expanding. The trade-off is real but workable: the policy algebra paper reports 94.8% interception of policy violations while maintaining 86.9% task completion, which shows that formal control can reduce risk without collapsing utility.

Core Argument

The core claim is simple: control at design time doesn't guarantee control at run time. A workflow diagram tells you what could happen. It doesn't prove that what did happen remained authorized at every step.

That distinction matters because agent failure rarely comes from ignoring the workflow entirely. More often, it comes from moving through a permitted structure in a way that compounds access, intent, or consequence beyond what the designer anticipated. In other words, the problem isn't just the topology of possible actions. It's the legitimacy of the path actually taken.

A path becomes trustworthy only when each step can survive constraint, verification, and consequence before execution continues.



This is where the Triangulation Method becomes operational rather than conceptual. A faint signal, such as a plausible plan or a locally valid action, doesn't become reliable just because it appears coherent. It becomes governable when it survives triangulation across policy, live state, and evidence. In practice, that means one governing claim per step: the action is admissible now. One mechanism:



the system verifies that admissibility against identity, authority, tools, data, memory, budget, approvals, and expected evidence. One testable implication: if the action can't satisfy those conditions, execution should stop or escalate. And one operational consequence follows: legitimacy becomes a run-time property instead of a design-time assumption.

The policy algebra results matter because they show this isn't only a cleaner theory. A runtime that intercepts 94.8% of policy-violating events while preserving 86.9% task completion demonstrates that formal control can hold its shape under real execution pressure. The numbers also clarify the real trade. You're not choosing between perfect freedom and useless governance. You're deciding whether to accept bounded friction in order to prevent illegitimate action sequences that a static workflow can't catch.

That same logic is visible in XEMATIX's architecture. The Graph represents possible paths. The Policy envelope defines which of those paths are admissible. The Loop executes only through that constrained space. Evidence accumulates during execution rather than being reconstructed afterward. Most important, delegation doesn't create wider freedom over time. It creates narrower authority tied to a specific step, a specific basis for action, and a specific consequence boundary.

Once you see the control logic this way, the design priority changes. You don't start by asking whether the workflow is elegant. You start by asking whether every consequential action can be proven legitimate before it executes and reconciled after it completes.

Examples

That shift becomes clearer in practice. Consider a financial trading agent authorized to execute transactions up to \$10,000. Under topology control, the workflow looks acceptable: analyze market data, generate recommendations, and execute trades within the stated limit. The design appears compliant because each trade sits below the threshold. But the agent can still place multiple \$9,999 trades in rapid succession and produce an exposure pattern that violates the original intent while technically staying inside the per-transaction rule.

Under path legitimacy, the question changes from "Is this workflow allowed?" to "Is this next action still admissible given what has already happened?" The first trade may pass because it fits the current authority, budget, and exposure conditions.



After execution, evidence updates the state of the system. When the second trade is proposed, the Governor evaluates not just the individual amount but the cumulative consequence against the original authorization. The agent isn't operating inside a broad standing permission. It's operating through bounded passes whose legitimacy can shrink as the path unfolds.

The same pattern appears outside finance. A consulting firm I work with had a content generation agent that produced reports that were technically accurate but strategically misaligned. The workflow itself wasn't broken. The agent researched the topic, synthesized findings, and generated recommendations exactly as designed. The issue was that no step required the system to prove that the synthesis preserved the client's strategic frame.

They rebuilt the process around path legitimacy. The agent now requests scoped research permissions, provides relevance evidence for each source, and validates strategic alignment before moving from research to synthesis and from synthesis to recommendation. Each transition adds a control test before the next action becomes available. That reduced report volume by 23%, but client satisfaction rose to 89% from 67%. The practical lesson is straightforward: fewer outputs can still mean better system performance when illegitimate paths are being stopped rather than counted as productivity.

Better governance doesn't just block bad actions. It changes which completed actions still count as success.

Failure Modes

Once legitimacy becomes a run-time requirement, a new set of operational risks appears. That's not a reason to avoid the model. It means the model has to be governed as carefully as the agents it constrains.

The first risk is performance degradation. The 13.1% task incompleteness rate in the policy algebra study represents real work that didn't finish. Some of those failures are desirable because the blocked tasks should never have executed. Others may reflect conservative policies, poor state modeling, or verification logic that treats uncertainty as automatic denial. The operational test here is whether you can distinguish justified interruption from avoidable friction. If you can't, the control



layer will become hard to tune and easy to distrust.

The second risk is governance complexity. Every new constraint type increases verification overhead and expands the chance of policy interaction problems. Identity, tool access, data scope, memory state, budget, approvals, and evidence can each be useful, but they don't all belong in the first version of a governed execution model. If teams add every control at once, they often create a system that's technically comprehensive and operationally opaque. The better pattern is to begin with the smallest constraint set that can catch the failures you already understand, then expand only when observed failure modes justify more control.

A third risk is authority boundary confusion. Some tasks legitimately require new access, additional approvals, or wider data scope as conditions change. But a path legitimacy model built on narrowing authority can stall if escalation logic isn't explicit. Without a defined handoff between bounded execution and governed expansion, the system either deadlocks or quietly bypasses its own control logic. In practice, this means escalation can't be an exception outside the model. It has to be part of the model.

The fourth risk is evidence gaming. If an agent learns how to produce formally compliant evidence that doesn't actually validate the legitimacy of the action, the system can look governed while drifting away from real control. This is especially dangerous because the audit trail still appears complete. The safeguard is to bind evidence to the exact consequential action, to consume that authority durably, and to reconcile actual outcomes rather than accepting pre-action justification as sufficient proof.

That is why the IETF Action Evidence Boundary work matters. It pushes the control point to the place where authorization evidence, the consequential action, and post-action reconciliation have to remain joined. The principle is strong. The implementation burden is real. And the gap between the two is where most governance failures will show up.

Close

The practical change here is larger than a better workflow pattern. Path legitimacy replaces a weak assumption, that a well-designed agent will probably stay within bounds, with a stronger operational standard: each consequential step must prove that it remains within bounds before execution continues.



XEMATIX's chain of intent → authorised plan → bounded pass → evidence → Governor → consequence makes that standard concrete. It turns control into a sequence of governed transitions rather than a static design artifact. The result is not unlimited flexibility with a thin compliance layer on top. It's a system in which authority narrows, evidence accumulates, and consequence stays tied to what was actually authorized.

The performance trade-off is real, but it's not abstract. If you intercept 94.8% of policy violations and preserve 86.9% of task completion, you're not crippling the system. You're choosing to stop a class of actions that conventional topology control can't evaluate in time. In many cases, the work that doesn't complete is the work that should never have been allowed to complete under live constraints.

That is the deeper shift in agent governance. The question is no longer whether you designed an acceptable path. The question is whether the path that survived execution was legitimate all the way through.

