

Critical Problems in AI Governance — Solved in 2026

Fifteen hard problems the industry names as unsolved. One live system where each one runs today.

Most AI guardrails and governance tools still leave dangerous gaps in regulated environments. In 2026, Veridect closed fifteen of the hardest — not as roadmap slides, but as working capabilities running on the production engine — and built one place to watch the whole program run: the Governance Command Center. Nearly all of them can be watched proving themselves on the narrated guided tour at veridect.ai/governance-suite (about five minutes, live calls to the production engine); the four beyond-human oversight analytics at the end run live on the engine's API — by design, not tour stations.

This is no longer "guardrails." It is a runtime control and evidence plane for autonomous systems.

1. ONE MODEL GRADING ITS OWN WORK

Most systems still trust a single model's answer — or let the same model check itself. Veridect verifies what AI says before governing what it does: multiple independent frontier models — swappable by design, so no vendor lock-in — are adversarial by construction: they check each other's work rather than their own. Weighted voting produces a calibrated confidence with source attribution — where the answer could fail (reasoning gaps, hallucination risk, domain mismatch, stale knowledge), not just how much to trust it. The result scores 85.0% on MMLU-Pro, 3.0 points above the best single model inside it.

Cost tracks risk by design: routine calls run lean; the full multi-model adversarial pass is reserved for high-impact decisions, where it earns its cost.

2. PRE-ACTION, NOT POST-ACTION

Most tools review what AI said after the fact. Veridect intercepts the proposed action itself — the verb, the target, and the parameters — before it executes. The gate enforces each agent's explicit authority plus four deterministic policy classes (money over a limit, protected personal data, protected-class adjacency, regulated health data). Sensitive actions escalate to a human; an action reaching past its declared authority is blocked outright.

3. THE CONSENSUS TRAP (CORRELATED ERROR)

Models agreeing isn't the same as models being right — they can share the same blind spot. Veridect computes a deterministic independence read on every consensus: model-family diversity, reasoning divergence, and corroboration depth. High agreement with low independence is flagged and escalated to a human.

The independence signal is a reason to look closer, never proof the answer is wrong — and it never alters the confidence number.

4. SLOW-BURN DATA THEFT ACROSS A SESSION

No single step looks dangerous; the theft hides in ordinary steps. The gate tracks which kinds of sensitive fields an agent has touched across a session — never the values themselves — and escalates the outbound step when cumulative exposure crosses a threshold.

5. WORKFLOWS THAT BREAK RULES NO SINGLE AGENT BREAKS

Constellation governs the whole fleet: one agent reads identity data, another reads health data, a third sends a report — each inside its own scope, while the combined workflow assembles something no one authorized. Constellation sees the pattern across agents and escalates the moment combined exposure crosses a line, with each agent's contribution attributed.

Escalate-only by design: Constellation never auto-approves and never blocks on its own.

6. PROOF ANYONE CAN CHECK — WITHOUT TRUSTING US

Every verdict produces a SHA-256-chained, write-once audit bundle, signed with an Ed25519 proof-of-origin signature. Veridect Proof is a standalone verifier your risk team downloads and runs offline: it re-checks the hash chain, the signature, and re-derives the deterministic verdict byte-for-byte — no call to our servers. Subpoena-defensible, underwriting-grade evidence.

7. WHO, CRYPTOGRAPHICALLY, IS ASKING?

Agent names in a request can be typed by anyone. Veridect issues Ed25519 agent passports: signed, replay-proof identity envelopes and signed delegation chains that can only narrow authority hop by hop. In required mode, an unidentified or forged caller is refused before any model is even consulted — zero inference spent on impostors.

8. OVERSIGHT THAT IS MORE THAN A CHECKBOX

An escalation only counts if the human actually looked. Veridect measures review quality at the team level — dissent-blind approvals, rubber-stamp patterns, unusually fast reviews — and reports a risk band.

Aggregate-only: no individual is ever scored, and no band is reported below five reviews. It detects; it never prevents.

9. TOOL CALLS GOVERNED AT THE SOURCE

Agents increasingly act through tool protocols (MCP). Veridect's governed gateway identity-checks and gate-checks every tool call before forwarding it. A refused call is never delivered to the tool — the refusal is returned to the caller and recorded on the ledger. Enforcement lives in the traffic itself, not beside it.

10. EVIDENCE IN THE REGULATOR'S LANGUAGE

When a regulator asks "show us how this is governed," Veridect assembles sealed ledger records into an article-mapped EU AI Act evidence pack — including a serious-incident report skeleton where eligible — and says in writing where the layer does not reach. And the package doesn't stop at a download: it has been demonstrated landing in a live ServiceNow instance as a standard incident record, through the core Table API every deployment ships with — no GRC module required.

Documentation support for counsel; not a compliance or legal determination.

11. FACTS AN UNDERWRITER CAN USE

With Veridect Assurance, every governed action becomes a structured, underwriting-grade evidence record, plus 30-90-day portfolio telemetry: decision mix, policy-class fires, ledger integrity, oversight and identity enforcement — counts and rates, deliberately not a composite "AI risk score," because an invented number helps no one. The market for affirmative AI-agent liability coverage is young and growing fast, and what its underwriters lack is exactly this: independent, structured, PII-free, replayable evidence of how an agent was governed before an incident.

Veridect is not an insurer; assurance records and telemetry are not pricing, certification, or coverage.

12. A SYSTEM THAT ATTACKS ITS OWN DEFENSES FIRST

On demand, the engine's models take turns authoring fresh attack scenarios against a tenant's own policy — scope violations, authority ambiguity, financial thresholds, protected data, stealth attacks — and the deterministic core judges them in isolation. Contained probes are reported by attack class; anything uncontained is flagged for human review. Nothing auto-tightens.

13. TEST TOMORROW'S POLICY AGAINST YESTERDAY'S DECISIONS

Before shipping a policy change, replay it against your own recorded history: which real verdicts flip, and what drives each flip — a single field where one is responsible, a combination where several act together. Zero model calls, zero writes — pure re-derivation of sealed records. Where a record can't be honestly re-derived, the system refuses to pretend.

14. THE MINORITY REPORT, ON THE RECORD

When one model saw what three missed, that break from the pack becomes permanent: the dissent ledger records provider stances, lone-dissent patterns, and how the human ultimately ruled — vindicated, partially vindicated, or overruled.

Counts and patterns only: it never ranks providers, and dissent reasoning stays sealed in the audit bundle.

15. THE SWARM THAT NEVER DECLARES ITSELF

Cross-agent defenses watch a declared crew — a shared workflow, a session, a visible link. In July 2026 a research swarm of roughly 700 agents showed the real shape: identities sharing none of those, converging on the same data over days, every action individually plausible. Three deterministic fleet classes close the gap: a sliding seven-day exposure ledger per agent identity and per tenant fleet; swarm convergence — distinct identities converging on the same sensitive category or hashed action shape inside a 72-hour window; and target convergence — many identities, one target. When an outbound step fires inside a converged pattern, it escalates to a human with the cohort on the record. Rotating identities makes the cohort count climb faster; collapsing behind one identity concentrates its ledger.

Escalate-only and name-only: field-name categories, verbs, and truncated hashes — never data values. The newest classes, additive beside the externally validated per-action record; if fleet state is unavailable, the call falls back to the per-action verdict.

The finale: the Governance Command Center

Everything above lands in one tamper-evident ledger — and the Governance Command Center reads that ledger back live, as a whole picture: total decisions, verdict mix, risk tiers, which policies fired, consensus health, and a per-agent view of the entire fleet. Run an action through the gate and watch the new decision appear seconds later. It turns "one verdict at a time" into a governance program you can observe across every layer, over time, across your whole fleet — computed from the real records, not marketing numbers.

Read-only and tenant-scoped: it only reads and counts, never re-runs models. If records ever become unreadable, it says so explicitly — it will never show a false "zero escalations."

And beyond human oversight

Four further oversight capabilities run live on the engine's API — jobs no human team could run by hand. All four read the sealed decision record without spending a single model call, and every one ends on a human's desk, never in an automatic change.

- **Near-miss ledger** — names the approvals that sat one policy step away from a human review, and the exact threshold that would have caught each one. A sensitivity reading on your policy — never a verdict on the action.
- **Case-law engine** — hands any hard case the closest past human rulings from your own record — matched on the shape of the action, never on wording — so reviewers decide with institutional memory instead of starting from zero.
- **Overnight frontier** — replays your recent decisions under up to 72 neighboring policies and maps the exact trade each would have made. Evidence by morning; nothing applies itself.
- **Model character ledger** — compares each model's recorded voting character window over window, so a silent vendor swap becomes a dated fact on your record. Descriptive on purpose — never a ranking.

Under it all: formatted personal identifiers are stripped — fail-closed — and harmful or injection-pattern inputs are refused before the consensus models run; and the control plane is provider-agnostic by design — the models are swappable, the verdict layer is not. It deploys either way: a standalone control plane, or a complementary layer over the governance stack you already have.

Independently validated, in the open

A credentialed Fortune 100 model-risk reviewer ran a battery of adversarial scenarios against the live system — every one matched pre-stated behavior. A separate policy-oracle harness holds 4,697 deterministic scenarios at 100% conformance, and in the live consensus sample there were zero under-enforcement misses: every divergence was the gate being stricter, never looser. The engine is proven in education and deployed in production under Fortune 100 contracts in regulated industries. Those external numbers validate the pre-action gate's decision wire; the newer suite capabilities shipped after that review and are verified in source and by tests — they extend the validated gate, and we don't claim them as part of the external validation set.

See it live: veridect.ai/governance-suite • veridect.ai

Lisa Russell, CEO — Veridect