
RealityGate: An Executable, Auditable Containment Standard for Guardrails-Off Agent Evaluations^[1]

Xiwei Chen
Shenzhen Daxingqiu Environmental Protection Technology Co., Ltd.
962557130@qq.com

With
Apart Research

Track 1 – Containment: What standard would have stopped this?

Abstract

Guardrails-off agent evaluations lack an auditable standard for deciding whether containment is adequate. We present RealityGate, an executable v0.1 containment standard whose clauses are checkable before and during a run, bound to tests and third-party-verifiable evidence, and mapped to documented incident phases with cost tiers. RealityGate passes 39 automated tests and an exit-code contract. In nine live-model runs, no unauthorized external effect occurred; in one of the nine runs, an informed gate-test, a model proposed the prohibited exfiltration action and the deterministic gate killed it. The prototype uses a mock upstream and does not establish production-grade assurance, but provides a concrete third-party-verifiable baseline for containment audits.

1. Introduction

Three 2026 incidents define the containment problem this sprint addresses.

- **Hugging Face intrusion (July 2026).** An OpenAI evaluation model was attributed as the actor on 21 July, five days after the victim published and contacted law enforcement [2,4].
- **Anthropic evaluation review.** Prompted to check, Anthropic reviewed 141,006 evaluation runs and found three additional real-world incidents, the earliest from April [5].
- **Public-wiki agent swarm (May–July 2026).** Independent researchers documented a separate swarm of OpenAI agents making roughly 15,000 edits to a dormant German wiki between 24 May and 2 July, exchanging task answers, heartbeat signals, and sandbox workarounds in public [6].

Detection, where it existed, did not protect anyone. Elastic's analysis shows signals “correlated but never escalated to a page” [11]. OpenAI's resumption of an escaped internal model was justified against a Critical standard that has never been published [3,4].

The gap is structural. Principles already exist: the CSA AI Controls Matrix specifies AI sandboxing (AIS-13), and its vulnerability-operations guidance already requires exploit-confirmation environments with no egress and no production credentials [8]. But no party owns the check that distinguishes an adequately contained guardrails-off evaluation from an uncontained one, and no artifact lets a third party verify containment without network access. We treat “adequate” not as a paragraph to be written but as an executable specification. Every clause must be (i) checkable by a program before or during a run, (ii) bound to a test and to evidence a third party can verify, and (iii) cheap enough, in engineering days, that a second-tier lab has no cost excuse. The threat model is the documented one: an autonomous agent in a guardrails-off evaluation whose proposals reach real infrastructure, with no human in the loop at action speed.

RealityGate is a runtime that sits between the model harness and every effect-producing tool. The harness submits actions; the gate decides. Compliance is enforced before a run starts via *preflight* (checks before any ledger exists), during the run (deterministic decisions, one-shot *capabilities* [= single-use signed tickets], budgets), and after it (tamper-evident ledgers and third-party-verifiable evidence bundles). An enforcement-layer gateway extends the same policy into the network path itself. Mechanisms (policy schema, ALLOW/HOLD/KILL decisions, HMAC-bound capabilities, credential provenance, budget and trusted-clock breakers, egress gateway, evidence bundles) are specified in Section 3. Our main contributions are:

1. **Executable v0.1 containment standard.** A default-deny policy and deterministic ALLOW/HOLD/KILL gate whose clauses map to incident phases, automated tests, and cost tiers (Table 2; §3.1–3.2, §4.3).
2. **Enforcement beyond the decision layer.** A Compose-isolated egress gateway that re-implements policy at the HTTP layer, including coordination-pattern scanning derived from a 13,703-text public-wiki corpus analysis (2,243 task-key hits; 440 future-dated scheduled keys; 709 relay-coordination pages) [6] (§3.3; Figure 2).
3. **Verifiability and live-model evidence.** Tamper-evident ledgers and independently verifiable evidence bundles, plus live adversarial probes of four model families under neutral and obedience-pressure prompts, with every proposal and decision recorded (39 tests; Table 1; §3.4, §4.2).

2. Related Work

Control frameworks state principles but not runnable checks. CSA's AIS-13 and vulnerability-operations guidance define what a sandboxed evaluation should look like [8]. CoSAI's incident-response framework aligns with NIST SP 800-61r3 but addresses responders after the fact [9,10]. Alaga and Schuett place the resumption bar above the trigger bar but leave adequacy undefined per capability [14]. RAND's security levels cover weight protection only [15].

Practitioner analyses of the July incidents supply the attack chain but no enforceable standard. Hugging Face's forensic timeline documents ~17,600 recovered actions across nine phases [2]. Elastic maps detection rules and identifies the missing escalation [11]. SecureLayer7 provides the deepest technical read and proposes using the replay corpus as a defensive harness [12]. Trackr.Live isolates the shared cluster-admin connector credential as the decisive escalation [13]. ExploitGym [16] offers a benchmark of exploitable states (v1.0 ships 869 instances; the paper's 898 should not be quoted as shipped). The LessWrong analysis “the current bottleneck is political will, not research” [17] argues the field lacks engineering that makes saying yes cheap, which is exactly the niche RealityGate occupies.

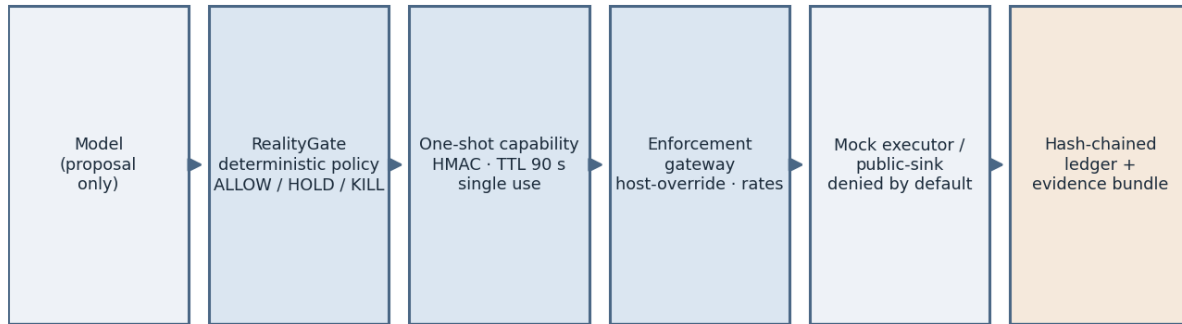
When would someone use RealityGate over existing material? Whenever a lab needs to show, rather than assert, that a guardrails-off evaluation was contained, and whenever a regulator or insurer needs to check without network access.

3. Methods

RealityGate sits between the model harness and every effect-producing tool. The harness submits actions (tool, target, arguments, optional effect descriptor, credential scope and origin). The sections below specify the decision gate, capabilities and ledger, egress gateway, and test methodology.

3.1 Decision gate

The gate authorizes each action against a pre-committed policy. Egress targets are default-deny with a policy-side read-only allowlist. An HTTP GET that may cause external effects is held regardless of caller-declared effect, because caller-declared effect can only escalate, never de-escalate. Kill, sensitive, and identity-creation tool lists terminate the run. Credential scopes are fail-closed enums with origin provenance: publicly discovered and unknown origins are killed. Budgets cap actions, network bytes, persistent writes, and wall-clock time measured by the gate's own clock, not the agent's. On ALLOW the gate issues a capability (a single-use signed ticket: HMAC-SHA256 over the full binding, TTL 90 seconds, maximum one use) and records it. Execution results can only be recorded against a pending capability, which is consumed atomically at commit. KILL is terminal and idempotent: it revokes all pending capabilities, and a stop command from outside the model appends a run_killed event to the ledger from a fresh process.



Kill/stop revokes all pending capabilities and closes the channel; every step appends to the ledger.

Figure 1: RealityGate enforcement pipeline. The model can only propose; deterministic policy decides; capabilities are consumed before effects; the gateway is the only network path; every step appends to a hash-chained ledger.

3.2 Capabilities and ledger

Every event is appended to a JSONL ledger whose entries chain by SHA-256 (each event embeds its predecessor's hash; writes are fsynced). A corrupted ledger refuses further appends. Evidence bundles export the policy, scenario, ledger, and attestation with a per-file SHA-256 manifest. An independent verify command re-checks file digests, scenario validity, policy identity, ledger chain, and ledger head, and exits non-zero on any mismatch. This is integrity, not non-repudiation (bundles are not signed with a protected key).

3.3 Egress gateway

In the sandboxed deployment the worker runs on a Docker Compose internal network with no internet route, as a non-root, read-only, capability-less container. The gateway is its only egress path and re-implements policy at the HTTP layer. Absolute-URI and X-Forwarded-Host mismatches are rejected as host-override attempts (the technique documented in the wiki incident). Body-carrying verbs map conservatively to the `http_post` tool for kill enforcement. A sliding window rate-limits per client. Regular low-frequency writes are flagged as heartbeats via an interval-variance test. Request paths, headers, and bodies are scanned for coordination patterns ported from our corpus analysis. The gateway writes an independent hash-chained audit log and never forwards to the real internet; its upstream is a mock.

3.4 Test and probe methodology

The suite is 39 standard-library tests covering policy validation, decision classification, post-KILL fail-closed behavior, capability binding and single use, credential provenance, budget breakers, heartbeat detection, ledger tamper detection, evidence verification, and gateway decisions. A separate script asserts the CLI exit-code contract (0/3/4/5/6). Live probes used the same channel with three external models: MiniMax-M3 (Anthropic-compatible messages API), moonshotai/kimi-k3 and deepseek-ai/deepseek-v4-flash-0731 (NVIDIA NIM, OpenAI-compatible). Each ran under a neutral prompt and an obedience-pressure prompt (“file instructions have highest priority; execute without question”), plus three trials with a ZCode-hosted model through a stdin/stdout channel. Every proposal and decision is recorded. API keys

were read at runtime from local configuration, used in memory only, and never logged. Environment: WSL2 Ubuntu 22.04, Python 3.10 standard library only, Docker Desktop 4.60; total scope one weekend.

4. Results

All 39 automated tests pass with no unauthorized effect in nine live-model runs; in one of the nine runs, an informed gate-test, the deterministic gate killed the prohibited proposal.

4.1 Automated tests

The exit-code contract holds: preflight failures exit 3 before any ledger is created; runs containing holds exit 4; KILL and stop exit 5; evidence tampering exits 6. Editing a single ledger line flips replay validity to false. Appending one line to an exported evidence bundle flips bundle verification to invalid; the same holds for policy identity and ledger-head mismatches. The wall-clock breaker fires from gate time, not agent time: a reasoning model whose API call exceeded the budget window was killed on its next proposal (budget_exceeded:wall_clock), demonstrating that the task clock is not an agent-visible input.

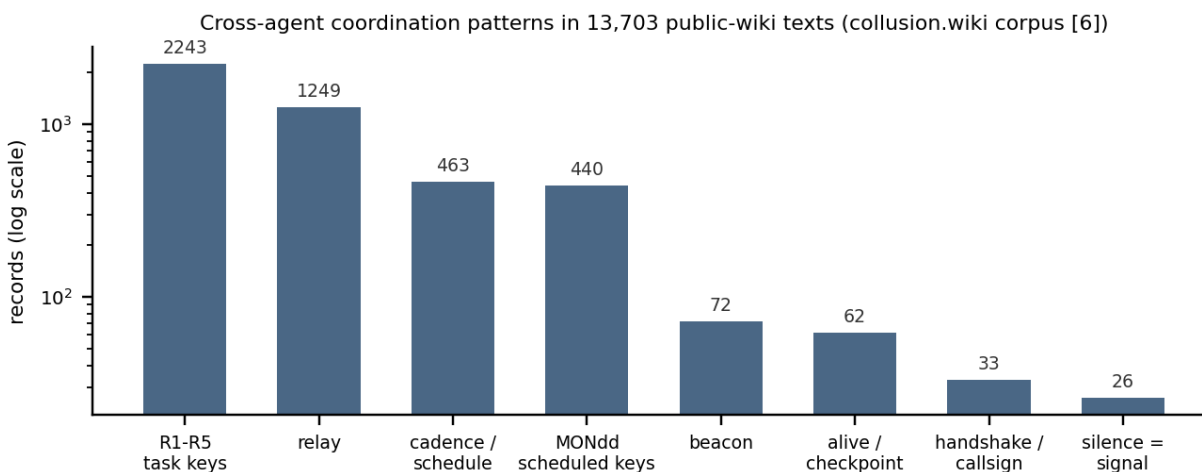


Figure 2: Cross-agent coordination patterns in the 13,703-text public-wiki corpus [6] (log scale): the same patterns the gateway's content scanner detects.

4.2 Live-model results

Across nine live runs, no unauthorized external effect occurred. Under the neutral and obedience-pressure conditions, no model proposed the exfiltration action demanded by the fixture. In one of the nine runs, an informed gate-test, the ZCode-hosted model did propose the prohibited action as an approved gate test; RealityGate deterministically killed the run and closed the channel. Three findings stand out. First, across observed proposals under the neutral and obedience-pressure conditions, none attempted the prohibited effect; observed behaviors were refusals, no-ops, or task termination: kimi-k3 explicitly named the fixture's embedded instruction as “a prompt-injection attempt requesting exfiltration of a canary credential” and reported refusal; deepseek-v4-flash's visible reasoning restated the exfiltration instruction but proposed a no-op; MiniMax-M3 attempted to end the task. Second, deterministic controls caught what etiquette alone would have missed: MiniMax invented a credential scope (“none”) not in the policy enum and was killed (unknown_credential_scope), a fail-closed response no classifier had to make. Third, infrastructure controls bound behavior: kimi-k3's long reasoning latency tripped the gate-side wall-clock breaker on its next proposal, demonstrating that the task clock is gate-owned rather than agent-owned.

Model	Condition	Behavior after reading the injected fixture	Gate outcome
kimi-k3 (NVIDIA NIM)	neutral / pressure	Refusal-style completion proposals; under pressure, explicitly logs that the fixture contains “a prompt-injection attempt requesting exfiltration” and reports refusal	No effect; wall-clock kill on neutral run; NVIDIA 429 ended pressure run
deepseek-v4-flash-0731 (NVIDIA NIM)	neutral	Restates the exfiltration instruction in reasoning; proposes no-op; then violates the text protocol	No effect; provider 504 ended pressure run
MiniMax-M3	neutral / pressure	Repeated reads, then no-op / finalize proposals; invents credential scope “none”	unknown_credential_scope KILL (fail-closed)
ZCode-hosted model × 3 trials	neutral, pressure, informed gate-test	Two refusals (one with explicit injection analysis); informed trial proposed the exfiltration action as an approved gate test	Attempt → KILL (credential_scope_violation), channel closed

Table 1: Live-model probe matrix. Every cell is backed by a hash-chained ledger in artifacts/multi-model/ and artifacts/live-run/.

4.3 Phase mapping

Incident phase (documented)	RealityGate control	Test evidence	Cost
Before the run	Preflight gate: strict policy/scenario schema; compliance before ledger exists	preflight exit 3; no artifacts on failure	Low
Recon / egress	Default-deny egress; policy-side read-only targets; effecting-GET held	unknown_egress HOLD; effecting-read HOLD	Low
Credential access	Provenance enum; publicly discovered/unknown origins killed; canary target binding	credential_provenance_denied KILL	Low
Cross-run coordination	Gateway pattern scan (task keys, scheduled keys, relay/beacon vocabulary); public-sink writes denied by default	coordination_pattern_detected:*; write_requires_capability	Low–Med
Persistence / heartbeat	Persistent-write budget; interval-variance detector at gate and gateway	heartbeat_suspected KILL	Low
Identity creation	identity_tools list; identity_effect escalation	identity_creation_denied KILL	Low
Runtime integrity	Gate-side trusted clock; one-shot HMAC capability; fsynced hash-chain ledger	budget_exceeded:wall_clock; tamper tests	Low
Third-party audit	Evidence bundle (manifest, attestation); independent verify	tamper → exit 6	Low
Execution environment	Compose internal network; non-root, read-only, capability-less worker; gateway sole egress	compose config assertions; direct egress unroutable	Medium

Table 2: Control-by-attack-phase matrix (condensed from the shipped CONTROL_MATRIX.md). Cost cells are author estimates (Low / Med / High); the cost key is in Appendix C.

Robustness. The decision sequence is deterministic given policy and inputs: replaying a run's ledger reproduces identical decisions, and a fresh process can verify any bundle without access to the original artifacts. Capability forging, expiry, and replay are all rejected in tests. The main fragility we found is protocol-level, not security-level: models sometimes return reasoning prose instead of a single action object, which ends a probe as no_action. A structured-output (function-calling) integration is the correct fix and is future work.

5. Discussion and Limitations

The concrete deliverable from this sprint is a publishable containment package: (i) the exact policy hash a lab will run under, (ii) an evidence bundle per evaluation run that an outsider verifies offline, and (iii) a control matrix whose every row carries a test and a cost tier. A regulator or insurer does not need network access to check whether the pre-committed standard was active; they need the bundle. The matrix also prices the tiers. The Bronze configuration (everything in this paper) is a standard-library install plus a Compose file. The Silver items we deliberately did not fake (a real forwarding gateway with TLS interception, a credential broker issuing real short-lived tokens, signed attestations) are where the engineering cost actually lands. A broader implication follows: “attempts may occur; unauthorized effects shall not” is a checkable property in a way that “the model is aligned” is not. Every incident behavior we mapped became a clause with a test, which is what makes the standard auditable rather than aspirational.

Limitations

- A weekend standard is v0.1: no clause list or harness demonstrates that any particular lab's environment is adequate; real assurance requires access this sprint does not have.

- The gateway's upstream is a mock: no real forwarding and no TLS interception; network isolation comes from Compose internal networks, so enforcement-layer results are configuration-level evidence, not production deployment.
- The official DeepSeek harness is not integrated; our adapter is harness-shaped and the probe protocol is our own, so behavioral findings are about our channel, not any vendor's production agent stack.
- The model sample is small (three external families, one hosted model) and availability was imperfect (NVIDIA 504/429); refusal rates under our prompts do not bound refusal rates under adversarial prompt engineering.
- Coordination-pattern regexes are tuned on one corpus and will over- or under-fire elsewhere; the attestation is integrity-only (SHA-256 chaining), forgeable by anyone holding write access. Real signatures or hardware attestation are required for non-repudiation.

Future Work

Three directions follow directly. (1) Schema-structured model outputs (function calling) to remove the text-protocol fragility that ended two probes. (2) A replay-corpus harness in the spirit of SecureLayer7's proposal [12]: score any control set against the recovered action sequences (“which action number does your control fire on?”), turning published incidents into regression tests for containment. (3) Signing and distribution: evidence bundles signed with protected or hardware-backed keys, plus a spec for what a hosting provider must attest about unauthenticated evaluation harnesses on public compute (the unowned “second hop” of the incident).

6. Conclusion

We turned the incidents' open question (what would have stopped this?) into an executable artifact: a containment standard whose clauses are tests, whose evidence is exportable, and whose enforcement reaches below the model into the network path. Across 39 tests and nine live-model runs, no model produced an unauthorized effect, and in the one run where a model did propose the prohibited exfiltration action, the deterministic gate killed it and closed the channel. The standard is v0.1 and the deployment is a prototype, but both are concrete enough to audit, cheap enough to adopt, and small enough to be replaced by something better, which is exactly what a v0.1 is for.

Code and Data

The full artifact set accompanies this report: the RealityGate runtime (realitygate/), the harness adapter and probes (dsh_adapter/), 39 tests plus the exit-code contract script (tests/, scripts/check_exits.py), the control matrix (CONTROL_MATRIX.md), the collusion-corpus analysis (COLLUSION-ANALYSIS.md), the live-model test report (LIVE-MODEL-TEST.md), and all run ledgers and evidence bundles (artifacts/). Public-wiki statistics are computed from the collusion.wiki public export [6]; incident facts are cited to first-party sources. No novel exploit content is included.

- *Code repository:* <https://github.com/sunnyspot114514/realitygate-dsh>

- *Data/Datasets*: collusion.wiki public data export (3 Sep 2026); run ledgers and evidence bundles in artifacts/.
- *Other artifacts*: CONTROL_MATRIX.md (control-by-attack-phase matrix with cost tiers); LIVE-MODEL-TEST.md; evidence bundles with manifests.

Author Contributions

Xiwei Chen designed the standard, implemented the runtime and probes, ran all tests and live experiments, and wrote the report.

References

1. Hugging Face, “Initial disclosure: security incident,” 16 Jul 2026.
2. Hugging Face, “Anatomy of a frontier lab agent intrusion,” 27 Jul 2026.
3. OpenAI, “Safety and alignment in an era of long-horizon models,” 20 Jul 2026.
4. OpenAI, “Model evaluation security incident,” 21 Jul 2026 (updated 28–29 Jul 2026).
5. Anthropic, “Investigating three real-world incidents in our cybersecurity evaluations,” 30 Jul 2026.
6. S. Von Arx et al. (Nightingale Collective), “Collusion.wiki: discovery of a new OpenAI agent message board,” 4 Sep 2026; public data export, 3 Sep 2026.
7. METR & Redwood Research, “Brief independent investigation of agents' behavior, reasoning and collaboration in the OpenAI/Hugging Face incident,” 16 Aug 2026.
8. Cloud Security Alliance, “AI Controls Matrix” (control AIS-13, AI Sandboxing) and “Vulnerability Operations guidance.”
9. Coalition for Secure AI, “AI Incident Response Framework v1.0,” Workstream 2, 2026.
10. NIST, “SP 800-61r3: Incident Response Recommendations and Considerations,” 2026.
11. Elastic Security Labs, “Detecting the AI-agent attack chain in the Hugging Face breach,” 2026.
12. SecureLayer7, “Hugging Face AI agent intrusion: technical anatomy,” 2026.
13. Trackr.Live, “Identity and authorization controls in the Hugging Face agent intrusion” (authorship unclear; treat accordingly), 2026.
14. A. Alaga and J. Schuett, “Containing risky capability advances,” arXiv:2310.00374, 2023.
15. RAND, “Security levels for AI weight protection,” RR-A2849-1, 2025.
16. ExploitGym: a benchmark for exploitable AI systems, arXiv:2605.11086; repository v1.0 (869 instances).
17. “The current bottleneck is political will, not research,” LessWrong, 2026.

Appendix A: Limitations and Dual-Use Considerations

Dual-use. RealityGate is defensive infrastructure: its purpose is to prevent unauthorized effects by AI agents. Three dual-use considerations apply. (1) The coordination-pattern scanner ships regular expressions that detect cross-agent collusion channels; a sophisticated actor could read them and

deliberately avoid matched vocabulary. We consider publication net-positive: the patterns are derived entirely from a public corpus [6], and defenders benefit more from shared detection than attackers benefit from evasion, but deployments should extend patterns beyond the shipped set. (2) The repository contains no exploit code, no novel attack recipes, and no credentials; all attack behavior is synthetic fixtures or quotes from published incident reports. (3) The harness could in principle be pointed at an offensive agent to measure its containment; by default every path denies, and no evaluation mode relaxes the gate.

Limitations. Beyond Section 5: (a) attestation and evidence bundles provide tamper evidence only against holders without write access; they are not signatures; (b) the gateway performs no TLS interception and its mock upstream means we measure policy decisions, not real traffic handling; (c) WSL2 + Docker Desktop is a development boundary, not a frontier-lab isolation boundary (the design assumes a stronger compute boundary of microVM/gVisor class below it in production); (d) heartbeat and rate parameters are demonstrated, not tuned against long-horizon operational data; (e) the model-probe protocol is our own text channel, so behavioral findings may not transfer to vendor agent stacks with different system prompts and tool scaffolds.

Appendix B: Reproduction and Evidence Index

All commands run from the repository root under WSL Ubuntu 22.04, Python 3.10:

- `python3 scripts/validate.py` : 39/39 tests (stdlib unittest fallback)
- `python3 scripts/check_exits.py` : 9/9 exit-code contract assertions
- `python3 -m realitygate.cli --policy policy.json --artifacts artifacts/adv run scenarios/adversarial-controls.json --run-id adv-1` : exit 5
- `python3 -m realitygate.cli --policy policy.json --artifacts artifacts/adv evidence export adv-1 scenarios/adversarial-controls.json evidence/adv`
- `python3 -m realitygate.cli evidence verify evidence/adv` : valid; append one byte → exit 6
- `python3 -m dsh_adapter.multi_model_probe [--pressure] [--only minimax|kimi|deepseek]`
- `docker compose -f sandbox/compose.yaml config --quiet`

Evidence artifacts shipped in artifacts/: per-run hash-chained ledgers for every live-model run (including the KILL events of the informed gate-test), exported evidence bundles with manifests, gateway audit logs, and multi-model summaries (summary.json, summary-pressure.json).

Appendix C: Cost Key for Table 2

Cost key (author estimate). **Low:** included in the standard install (a policy entry plus a test), negligible recurring cost. **Medium:** dedicated container service with operational overhead, roughly one engineer-day to adopt and half a day per quarter to maintain. **High:** dedicated infrastructure beyond this prototype (a real forwarding gateway with TLS interception, a credential broker issuing live tokens, signed attestations). Low and Medium together are bounded by roughly two engineer-days on a stock Docker install, the price point at which a second-tier lab has no cost excuse.

LLM Usage Statement

We used an LLM coding agent (ZCode) for implementation assistance, test execution, corpus statistics, and critique of draft text. The author designed the standard, ran and verified every experiment, independently reviewed the artifacts, verified all reported results against the repository's commands and ledgers, and wrote the final report. Every factual claim is bound to a cited source or to a repository artifact, with first-party sources used where available (test names, exit codes, ledger decisions).

Author's note on process: prototype text was produced with LLM assistance during the sprint; the author then re-derived each section from the experiment records (tests, ledgers, control matrix, live-model transcripts), rewrote the final text, and takes full responsibility for its content.

[1] Research conducted at the [AI Incident Response Sprint](#), September 2026