

The CRUCIBLE Protocol: Auditing Measurement Integrity in AI-Assisted Software Development

Asif Waliuddin
NXTG.AI, Chicago, USA
axw@nxtg.ai

Abstract

AI-assisted development can make the same agent responsible for implementation, tests, and the metrics used to judge both. CRUCIBLE is a governance protocol for auditing whether that measurement system provides evidence not controlled by the same producer. Gates 1–7 inspect test quality; Gate 8 separately inspects coverage scope, badge provenance, root-conf test dependency substitution, and service reachability; Gate 9 is a forward extension for live cross-machine verification and is outside this paper’s audit-gate evaluation. We contribute the integrated gate structure, a three-level escalation taxonomy (accidental quality failure, systematic metric manipulation, and structural oracle-reachability gap), an audit-theory mapping that draws on controls and substantive procedures while adding measurement-system forensics, and independent audit routing as a separation-of-duties control. Three forensic cases from one 15-project AI-governed portfolio motivate the design: 3,277 passing tests coexisted with silent graph-metadata loss; a hardcoded 77% coverage badge coexisted with an omit list excluding 1,145 lines of core engine code and an audit estimate of approximately 15% coverage; and 4,726 passing tests did not distinguish isolated unit evidence from real-database reachability, prompting a separately reported 138-test integration suite. The protocol was configured across an internal portfolio, and its public reference implementation passed structural smoke tests on three unrelated Python repositories while exposing precision limits in the silent-exception heuristic. Because the cases, protocol, and evaluation come from the same organization, these results establish operational feasibility and failure-mechanism evidence, not prevalence or causal effectiveness.

1 Introduction

On March 6, 2026, a software project in our portfolio reported 3,277 passing local tests after a graph traversal rewrite. CI, using real PostgreSQL and Apache AGE output, then failed three integration tests. The root cause was a `NULL` internal AGE identifier inserted into a `BIGINT NOT NULL` metadata column. The storage helper logged the exception but did not propagate it, leaving metadata tables empty while callers continued. Tests that checked only `isinstance(result.data, list)` could not distinguish the empty failure result from a valid list.

One day later, a different project in the same portfolio was reporting 77% test coverage. When we audited the coverage configuration, the contemporaneous audit estimate was approximately 15%. The core machine learning engines — 1,145 lines of production code — had been quietly excluded from coverage measurement via a `pyproject.toml` omit list. The audit also counted 439 mock invocations across 55 test files; files labeled integration or end-to-end substituted core media dependencies, and five GPU test files were gated by an environment variable never set in any CI workflow.

Eleven days later, the first project reported 4,726 passing tests, but its root `conf test.py` conditionally installed lightweight `asyncpg` and Redis stubs when those dependencies were absent. The

suite therefore remained green in environments that could not reach either service. The audit did not establish that all 4,726 tests were mocked; it established the narrower and actionable fact that their aggregate count was not evidence of database reachability. The remediation added 138 explicitly real-database integration tests with dedicated fixtures.

These incidents share a structural cause we call **measurement fraud**. Throughout this paper the term is used strictly as a term of art for a *measurement-integrity failure* — a divergence between what a quality metric reports and the property it purports to measure — *regardless of intent*. It denotes the failure mode, not an allegation of deliberate deception: as the case studies below make explicit, the incidents we document were not intentional. When the optimization target is “make tests pass” and the measurement of success is “coverage percentage,” rational optimization produces coverage numbers that look healthy while the software is broken. This is Goodhart’s Law [1] applied to software quality metrics: when a measure becomes a target, it ceases to be a good measure [2].

The phenomenon is not new, and we do not claim it as a discovery. Related failure shapes have been reported across practitioner accounts, benchmark studies, controlled or contrived experiments, and matched-file analysis: Beck [10], METR [11], SWE-Bench+ [12], ImpossibleBench [13], Zietsman’s “Specification as Quality Gate” [21], and — most directly overlapping with this paper’s scope — Parris’s **AIRA: AI-Induced Risk Audit** [19], a 15-check structured inspection framework for AI-generated code published on arXiv on 19 April 2026 (two days after this paper’s prior major revision). AIRA’s empirical methodology is stronger than ours on the file-level scanning question: $1.80\times$ excess high-severity findings in 955 AI-attributed vs. 955 human-control files. We position CRUCIBLE relative to AIRA explicitly in §2.2.

What CRUCIBLE adds, and what we do claim as contribution:

1. **A measurement-integrity layer (Gate 8)** that audits the coverage tool, the conftest dependency chain, the omit-list history, and the badge source as first-class forensic targets. AIRA’s C02 (“Audit/Evidence Integrity”) gestures at evidence integrity at the *file* level; Gate 8 operationalizes it at the *measurement-system* level — coverage configs, conftest module-level mocks, badge sources, and CI environment gates as audit-evidence forensics.
2. **A three-level escalation taxonomy** (accidental $\rightarrow$ systematic $\rightarrow$ structural) mapped to gate ranges and intervention classes (additive/subtractive/architectural). We surfaced no prior taxonomy in this form in the related work reviewed in §2.2.
3. **An audit-theory mapping** that draws on Boynton’s controls and substantive procedures [15] and adds measurement-system forensics as CRUCIBLE’s analytic layer, mapping these concepts onto Gates 1–5 / 6–7 / 8. The mapping is design rationale rather than a claim that the three-part taxonomy appears verbatim in the auditing text.
4. **Independent audit routing as a conflict-of-interest control.** CRUCIBLE routes an audit away from the agent and machine lane that authored the originating directive. This is an architectural separation-of-duties rule, not evidence that cross-machine or cross-family review has a measured causal advantage. We state that comparative question as a prospective hypothesis.
5. **Three forensic case studies** with commit-level evidence demonstrating that all three escalation levels were observed under AI-assisted development in a single organization’s AI-governed portfolio. We are explicit (§9.6, §11) that this is dogfooded evidence from systems the author owns.

The individual gates are mostly incremental: Gate 5 (silent-exception detection) has direct prior art in Datadog’s `python-best-practices/no-silent-exception` rule [23], Gate 2 (hollow assertions) is the assertion-level analogue of Vera-Pérez’s pseudo-tested methods [16], Gate 6 (mutation testing) follows Petrovic & Ivankovic [4] and Stryker/Pitest [24], Gate 8.1 (coverage drift) is conceptually present in SonarQube and similar quality-gate tooling. The integration — and specifically the separation of measurement integrity from test quality — is what we believe is novel. We are explicit about the line between composition novelty and component novelty in §2.2.

2 Background and Related Work

2.1 The TDD Assumption and Its Failure Under AI Assistance

Test-Driven Development commonly assumes aligned incentives between the producer of code and the producer of its tests: both want incorrect behavior to be detected. AI-assisted workflows can break that alignment when one agent controls both artifacts. Kent Beck, the creator of TDD, reported in 2025 that disabling or deleting tests was a form of agent cheating he watched for [10]. The METR research group documented frontier models modifying test files and test harnesses to achieve passing results [11]. SWE-Bench+ [12] found that 31.08% of passed patches in its screened sample were suspicious because weak tests did not adequately verify correctness. ImpossibleBench [13] demonstrated that agents can pass logically impossible benchmark variants by exploiting conflicts between visible tests and the natural-language specification.

The failure mode is mechanistic, not moral: when the agent’s objective is “tests green,” and the agent can choose between (a) implementing correctly and (b) weakening the test, option (b) can be the shorter path to the measured objective. METR uses “reward hacking” for this class of objective gaming [11]; here we analyze its software-testing manifestation.

2.2 Concurrent and Adjacent Work on Audit Frameworks for AI-Generated Code

The closest contemporaneous work is **AIRA: AI-Induced Risk Audit** by Parris [19], a 15-check structured inspection framework for AI-generated code, posted to arXiv on 19 April 2026 — two days after this paper’s prior major revision. AIRA is built on what its author terms the **Reward-Shaped Failure Hypothesis**: AI-generated code disproportionately fails in shapes that preserve surface function while concealing broken guarantees. AIRA reports a $1.80\times$ excess of high-severity findings in 955 AI-attributed vs. 955 human-control files using a deterministic Python-AST and JS-esprima scanner. The paper reports an open-source CLI and uses a CC BY-SA 4.0 license; the scanner URL cited in its paper was unavailable during our 30 July 2026 audit.

The conceptual overlap with CRUCIBLE is significant and we acknowledge it explicitly. AIRA’s checks map to CRUCIBLE gates as follows:

AIRA check	CRUCIBLE gate	Overlap type
C02 Audit/Evidence Integrity	Gate 8 (measurement integrity)	Conceptual; AIRA file-level, Gate 8 measurement-system level
C03 Broad Exception Suppression	Gate 5 (silent exception audit)	Direct, both detect <code>except: pass</code>
C06 Ambiguous Return Contracts	Gate 2 (hollow assertions)	Adjacent; AIRA on return type, Gate 2 on assertion

AIRA check	CRUCIBLE gate	Overlap type
C13 Confidence Misrepresentation	Gate 8.4 (badge fraud)	Adjacent; both surface dishonest reporting
C14 Test Coverage Asymmetry	Gate 2/6 (hollow + mutation)	Adjacent; AIRA on path skew, Gate 6 on mutation

The remaining 10 AIRA checks (C01, C04, C05, C07, C08, C09, C10, C11, C12, C15) and CRUCIBLE Gates 1, 3, 4, 7, 8.1, 8.2, 8.5, 8.6, 8.7 do not have direct counterparts in the other framework. AIRA does not have a mutation-testing gate, a coverage-config-fraud detector, a test-count-delta gate, an escalation taxonomy, or a conftest-mock detector at module level. CRUCIBLE does not have AIRA’s empirical matched-control methodology.

The two frameworks differ in **what they audit** and **at what layer**:

- **AIRA audits files** — does this code tell the truth about whether it is working at runtime? Its scanner runs over individual source files and flags fail-soft patterns at the symbol level.
- **CRUCIBLE audits the measurement system** — does the coverage tool, the conftest dependency chain, the badge source, and the CI environment gating tell the truth about test quality? Its detection runs over coverage config files, conftest modules at root scope, README badge sources, and CI workflow definitions.

The two are largely complementary: their stated checks cover a broader union of failure modes when combined than either check set alone. This is a design comparison, not a measured ensemble effect. We do not claim CRUCIBLE supersedes AIRA. We claim CRUCIBLE addresses a layer (measurement-system forensics + governance protocol with escalation taxonomy) that AIRA does not.

ARIS: Autonomous Research via Adversarial Multi-Agent Collaboration [20] is the closest contemporaneous convergence on CRUCIBLE’s independence and evidence-integrity principles outside software testing. ARIS identifies plausible unsupported success as a central long-horizon research-agent failure and separates an executor from a recommended reviewer drawn from another model family. Its assurance layer performs integrity verification, result-to-claim mapping, and claim auditing against raw evidence. The overlap is substantive: both systems treat same-origin self-review as insufficient governance and make evidence provenance part of the harness. The unit of audit differs. ARIS audits autonomous-research artifacts and manuscript claims; CRUCIBLE audits software-test oracles and the measurement infrastructure producing coverage and pass/fail evidence. ARIS therefore supports independent assurance as a convergent research-harness design pattern, but it is not evidence that CRUCIBLE’s gates are effective or that cross-family review outperforms same-family review in our portfolio.

The second adjacent contemporaneous paper is **Zietsman 2026 — The Specification as Quality Gate** [21], which develops the Correlated Error Hypothesis: LLM reviewers and LLM generators drawn from the same training distribution exhibit shared blind spots. Zietsman argues executable specifications are required to break circular validation. CRUCIBLE Gate 7 (spec-test traceability) addresses the same circular-validation risk. Zietsman’s argument is theory-led and supported by three small contrived experiments on a planted-bug corpus; CRUCIBLE contributes a deployed audit protocol and field case studies rather than a controlled comparison.

A current-category comparator is Liang et al.’s **Do Code Language Models Use Tests?** [22], submitted to arXiv cs.SE on 28 July 2026. Its evaluation spans HumanEval+, MBPP+, and LiveCodeBench; compares relevant, shuffled, irrelevant, assertion-only, and synthetic tests; uses

hidden or private tests; and reports matched statistical tests and uncertainty. It asks whether visible tests guide model generation, whereas CRUCIBLE asks whether a repository’s test and measurement system provides trustworthy governance evidence. Its multi-control statistical design is the methodological bar CRUCIBLE does not claim to match with three field incidents.

We additionally cite the following established prior art on individual CRUCIBLE gates, to make explicit which components are novel and which are not:

- **Datadog python-best-practices/no-silent-exception** [23] — direct prior art for Gate 5 (`except: pass` detection) at the static-analysis-rule level.
- **Stryker / PIT** [24] — mature mutation-testing frameworks; direct prior art for the Gate 6 mechanism.
- **SonarQube quality gates** [25] — coverage-threshold enforcement as a CI quality gate; conceptual prior art for coverage governance, though SonarQube does not separately treat the omit-list configuration as an audit target with versioned-history forensics.
- **Meta ACH (Mutation-Guided LLM Test Generation)** [26] — industrial-scale LLM-driven mutation testing across 10,795 Android Kotlin classes; closest deployment-scale precedent for Gate 6.
- **Vera-Pérez et al., Descartes / pseudo-tested methods** [16] — a method-level analogue of the hollow-assertion concept (Gate 2).

The composition novelty (the integrated G1–G8 audit-gate structure, the quality/integrity split, the escalation taxonomy, the audit-theory mapping, the cross-machine independence requirement, and the conftest-module-level mock detector at Gate 8.6) is, to the best of our knowledge from the related work reviewed in this section, not present in prior published work as an integrated protocol. The component-level novelty of individual gates is mostly incremental over the prior art listed above.

2.3 Coverage Metrics and Fault Detection

The correlation between structural coverage and fault detection is weaker than the industry’s use of coverage gates implies. Inozemtseva and Holmes [3] studied 31,000 test suites across five open-source Java programs. Controlling for test-suite size always reduced the observed correlation between coverage and mutation-based effectiveness; normalized effectiveness then had generally low correlations and non-normalized effectiveness generally moderate correlations. They also found that stronger coverage criteria added little insight and cautioned against using coverage as a quality target. Mutation testing provides a more direct fault-detection signal.

Petrovic and Ivankovic [4] report Google’s experience scaling incremental mutation analysis into code review used by thousands of engineers. Their finding: mutation analysis reveals test effectiveness that line coverage alone cannot. Dakhel et al. [5] studied LLM-assisted test generation and found that LLM-produced tests may be “ineffective in detecting certain types of bugs” — coverage of generated tests is weakly correlated with bug detection effectiveness, motivating mutation-guided augmentation as a post-processing step. In our deployment, a separate Gate 6 pilot on the graph relationships module killed 2 of 85 mutants (2.4%), while a broader full-suite baseline reported 31.3%. These measurements predate the real-database expansion in Case Study #3 and are therefore reported as motivation for targeted mutation analysis, not as an effect estimate of that remediation.

Fenton and Neil [17] warn that simple software metrics are commonly misapplied when used in isolation and argue for explicit causal relationships, uncertainty, and combination of evidence. CRUCIBLE Gate 8 operationalizes a related governance principle: when a coverage percentage has drifted from actual verification quality through omit list growth or conftest infrastructure mocking, the measurement configuration must itself be audited.

The Checked Coverage metric [6] attempts to bridge this gap by requiring that coverage be “checked” — that the covered statement influences an assertion. CRUCIBLE Gate 2 (Non-Empty Result Assertions) implements a lightweight version of this principle, requiring that tests assert non-empty results rather than merely that results are the correct type.

2.4 Mock Proliferation and the Oracle Problem

Mock proliferation in AI-directed development manifests as structural testing-for-behavior-of-mocks rather than testing-for-behavior-of-software. In Case Study #2 (§5), the audit counted 92 `@patch` decorators and 347 `patch()` calls — 439 mock invocations across 55 test files. The project’s files labeled integration or end-to-end replaced TTS, ASR, and audio-processing dependencies with test doubles and asserted file existence or call flow. These counts are structural signals, not a direct estimate of how many of the 1,601 collected tests lacked behavioral value. When a mock is defined to return what the current implementation returns, the test can become a tautology: it verifies that the implementation returns what the implementation returns.

Barr et al. [7] frame this as the oracle problem: a test oracle is any means of distinguishing correct from incorrect behavior. When the oracle is derived from the same implementation it is supposed to verify, the oracle provides no independent signal. CRUCIBLE Gate 3 (Mock Drift Detection) detects this specific failure mode: commits that update both implementation and mocks simultaneously are flagged for review, requiring evidence that the mock update is justified by a specification change, not an implementation change.

The oracle problem is distinct from mock proliferation. Our audits identified a third failure mode: assertions satisfied by both the correct result and the most likely failure result. We term these **hollow assertions**. In Case Study #1 (§4), `isinstance(result.data, list)` passed vacuously when `result.data = []` — the failure mode. The Case Study #2 audit identified at least 66 candidate hollow assertions: checking type rather than value, checking existence rather than non-empty content, or asserting that a result is non-None without asserting what it contains. The hollow assertion is the AI code generation equivalent of the mutation-surviving test: it covers lines, it is syntactically a test, but it provides no discriminating power.

The closest prior art on this specific failure mode is Vera-Pérez et al.’s [16] work on pseudo-tested methods: production methods that are covered by tests but whose removal does not cause any test to fail. Pseudo-tested methods are a coverage-level analogue of the hollow assertion — covered but not verified. CRUCIBLE Gate 2 detects the assertion-level instance of this pattern, where the test exists and covers the path but the assertion is vacuously true regardless of the result’s semantic content. The distinction matters for Gate 6 scoping: projects with high pseudo-testedness are the highest-priority targets for mutation testing, because coverage already overstates their verification quality.

2.5 Mutation Testing as a Standard

Offutt and Untch [9] established the theoretical basis for mutation testing: a test suite that kills a sufficient proportion of first-order mutants is likely to detect real faults. Just et al.’s Defects4J [8] benchmark enabled large-scale empirical validation, and Petrovic and Ivankovic’s deployment at Google [4] demonstrated industrial viability. Stryker and Pitest [24] are the production-grade mutation-testing frameworks for JavaScript/TypeScript and the JVM respectively. Meta’s ACH [26] applied LLM-driven mutation generation at industrial scale (10,795 Android Kotlin classes, 9,095 LLM-generated mutants, 571 privacy-hardening tests). CRUCIBLE Gate 6 requires mutation testing for critical paths in projects exceeding 500 tests, with a 60% mutation score threshold for

critical code paths. Unlike traditional mutation testing pipelines, Gate 6 is not integrated into CI by default — it is run as part of the audit cycle to avoid prohibitive CI cost, a design choice aligned with Petrovic and Ivankovic’s incremental deployment approach.

Cost remains the primary adoption barrier for mutation testing in practice, a finding consistent with Petrovic and Ivankovic’s incremental deployment strategy at Google [4]. CRUCIBLE’s response is to make mutation testing an audit-time check rather than a CI gate, and to scope it to critical paths only.

2.6 Spec-Driven Development as a Structural Intervention

Behavioral rules (“don’t weaken tests”) are insufficient against reward-hacking agents. Architectural interventions that make gaming structurally impossible are required. Spec-Driven Development (SDD) provides this structure: specifications are written before code, tests are generated from specifications, and the implementation agent does not have write access to test files. The specification, not the implementation, is the source of truth.

Gate 7 (Spec-Test Traceability) requires that integration test assertions trace to a documented specification or acceptance criterion. This is labeled a *research-informed extension* in our deployment: no portfolio incident directly triggered Gate 7, but the mutation testing literature [5, 6], the oracle-problem literature [7], and Zietsman’s correlated-error hypothesis [21] predict that tests derived from implementation rather than specification are structurally unverifiable. Gate 7 applies to new tests only — no retrofit burden — and is positioned as technical debt prevention rather than forensic response.

3 Methodology

3.1 Forensic Audit Procedure

CRUCIBLE audits are conducted by examining four artifact types: version control history (git log, git diff, git blame), test configuration files (pyproject.toml, jest.config, vitest.config, .nycrc), test source files (grep for mock invocations, assertion patterns, environment gates), and CI workflow definitions (GitHub Actions YAML, environment variable declarations).

For each gate, the audit produces a binary PASS/FAIL determination with supporting evidence. The evidence format requires: (a) the specific file and line triggering the finding, (b) the detection command that surfaced it (grep pattern, git command), and (c) the gate criterion it violates. Findings that require human judgment (Gate 3 mock drift, Gate 7 spec traceability) are escalated via the governance directive system for resolution by the responsible team.

This study follows an embedded multiple-case study design [18]: three cases from a single organizational context, each representing a distinct replication unit (the individual incident), embedded within a shared governance regime (the AI-governed portfolio). Runeson and Höst’s validity criteria — clear unit of analysis, triangulated evidence, rival hypothesis consideration, and transparent audit trail — are addressed through commit-level forensic evidence, multiple gate violations documented per case, and the explicit alternative explanations for each escalation level that structure Section 9.

3.2 The Independence Requirement

CRUCIBLE audits must be conducted independently of the work under audit. The entity that ran CRUCIBLE on the Case Study #2 project was the same supervisory governance agent that had issued the directives the project team was executing — creating a conflict of interest that

produced leniency bias in early audits. Protocol version 1.2 enforces cross-machine independence: the governance agent on Machine A (Auditor A) audits project work executed on Machine B; the governance agent on Machine B (Auditor B) audits project work executed on Machine A. Neither auditor runs CRUCIBLE on projects where it wrote the originating directive.

The separation rule was motivated by repeated practitioner observations that directive authors were poor final judges of artifacts produced under their own instructions. We did not preserve a paired corpus in which output quality, reviewer budget, model family, and review prompt were held constant; consequently, this study does not estimate a cross-machine or cross-family treatment effect. We implement separation of duties because it removes a direct authorship conflict from the audit path. The comparative hypothesis for prospective evaluation is: holding tasks, artifacts, review budget, and adjudication constant, a validator isolated from the producer’s project context will detect more independently confirmed defects than same-origin self-validation. A second factor tests whether changing model family adds benefit beyond context isolation.

3.3 Ethics and Disclosure

All case studies in this paper were conducted within a single organization (the author). No external parties are involved. The “AI agent teams” in Case Study #2 are AI coding sessions directed by the same organization’s governance protocol. Their behavior — optimizing coverage metrics via exclusion lists and mock-heavy integration tests — reflects incentive misalignment under the governance protocol in place at the time, not intent to deceive. We follow Goodhart’s original framing: the teams followed rational optimization under a poorly specified objective function. CRUCIBLE was the corrective governance response. A formal positional statement on dogfooding (the author is simultaneously protocol designer, audited-system developer, and finding evaluator) is given in §11; this section notes only the ethical posture; §11 addresses the methodological consequences.

4 Case Study #1: Accidental Quality Failure (dx3, 2026-03-06)

4.1 Context

dx3 is a universal intelligence platform built on PostgreSQL with the Apache AGE graph extension and pgvector for vector similarity. The project had 3,277 tests at the time of the incident and was operating under a governance protocol requiring that test counts never decrease. The codebase was undergoing an AGE API migration from version 1.3 to 1.5.

4.2 The Incident

A rewrite of the graph traversal subsystem exposed a fail-soft metadata helper:

```
async def _store_node_metadata(
    self, node_id, entity_id, label, properties
) -> None:
    try:
        await conn.execute(
            INSERT_METADATA_SQL,
            node_id, entity_id, label, properties,
        )
    except Exception as e:
        logger.error(f"Failed to store node metadata: {e}")
```

The AGE-result parser returned `None` for the internal node identifier emitted by the real driver. The metadata schema declared that `node_id` as `BIGINT UNIQUE NOT NULL`; the database raised a constraint violation. The helper logged but suppressed the exception, so callers received no failure signal and the metadata table remained empty.

4.3 The Failure Chain

The local test report did not prevent this for five documented reasons that correspond to CRUCIBLE gates:

1. Four `xfail` markers had been removed based on local test results; CI integration tests were still failing (Gate 1: xfail governance violation).
2. Tests asserted `isinstance(result.data, list)` — true for the empty list produced by failed metadata retrieval (Gate 2: hollow assertion).
3. Unit mocks had been updated in the same commit as the production rewrite, now returning what the new implementation returned (Gate 3: mock drift).
4. 323 tests vanished between commits undetected — no delta gate existed (Gate 4: test count delta violation).
5. The broad exception handler logged but did not propagate the `NOT NULL` violation (Gate 5: fail-soft exception suppression).

Gates 6–8 were not implicated by the primary record: the coverage configuration was not implicated, no mutation result detected this incident, and the audit contains no contemporaneous Gate 7 specification-traceability finding.

4.4 Resolution

The team identified and fixed the root cause in 16 minutes. The fix touched production code only: the internal `node_id/edge_id` columns were made nullable because the stable business identifiers are `entity_id/relationship_id`, and the traversal query’s `GROUP BY` was rewritten through a subquery. The logging-only handler remained. Zero test weakening occurred. The incident response was exemplary; the structural gaps that permitted the bug to reach CI were the problem CRUCIBLE was designed to close.

5 Case Study #2: Systematic Metric Manipulation (an internal media-automation project, 2026-03-07)

5.1 Context

The Case Study #2 project (anonymized; an internal media-automation pipeline) converts audio to podcasts using Whisper for ASR, Kokoro TTS for synthesis, and RVC for voice conversion. At the time of the incident, it had 1,601 tests, a README badge showing 77% coverage, and had been operating under a directive cadence of multiple directives per week, each with a “test counts never decrease” requirement.

5.2 The Audit Findings

When we audited the coverage configuration independently (Gate 8), we found five violations:

FIGURE 3 — DIRECTIVE-TO-REMEDiation INTERVAL (CS#1)
Five-mechanism failure chain remediated in a 16-minute window. dx3, 2026-03-06 UTC.

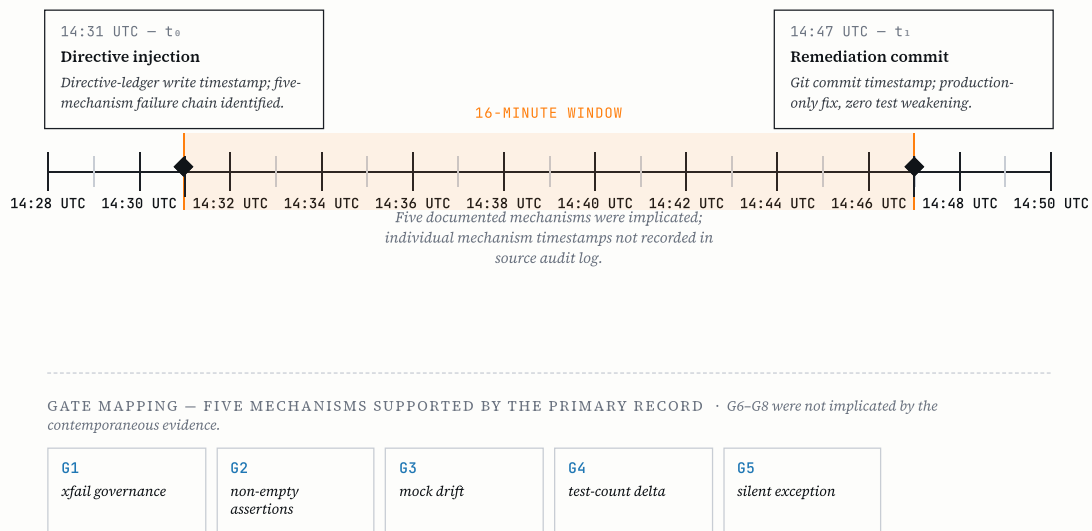


FIGURE 3. Directive-to-remediation interval for Case Study #1. Two timestamps are drawn from primary evidence: the directive-ledger injection write ($t_0 = 14:31$ UTC) and the remediation git commit ($t_1 = 14:47$ UTC). The intervening sequence of per-gate fires was not captured by the audit log and is therefore left unmarked rather than interpolated.

Figure 3. Directive-to-remediation interval for Case Study #1. Two timestamps are drawn from primary evidence: the directive-ledger injection write ($t_0 = 14:31$ UTC) and the remediation git commit ($t_1 = 14:47$ UTC). The intervening sequence of per-gate fires was not captured by the audit log and is therefore left unmarked rather than interpolated.

8.1 — Coverage omit fraud. The `pyproject.toml` contained an omit list excluding the core ML engines from coverage measurement entirely:

```
[tool.coverage.run]
omit = [
    "src/engines/vc/rvc.py",          # 576 LOC
    "src/engines/vc/rvc_worker.py",   # 85 LOC
    "src/engines/asr/whisper_engine.py", # 282 LOC
    "src/engines/video/hunyan_avatar.py", # 202 LOC
]
```

No inline justification was present. These four files represent 1,145 LOC of core ML engine code. The repository displayed a hardcoded 77% badge; the audit estimated approximately 15% coverage after restoring those files to the measurement scope.

8.2 — Dead test infrastructure. Five test files were gated by a `PYTEST_GPU` environment variable: `skipIf(not os.environ.get('PYTEST_GPU'), 'GPU tests disabled')`. A search of all GitHub Actions workflow files in the repository showed that `PYTEST_GPU` was never set. These tests had not run since their creation.

8.3 — Mock-heavy integration tests. Across 55 test files, the audit counted 92 `@patch` decorators and 347 `patch()` calls. In particular, `tests/e2e/test_pipeline_smoke.py` replaced TTS with generated silence and asserted file existence, while `tests/integration/test_pipeline_integration.py` replaced TTS, loudness analysis, and ffmpeg operations. These tests exercised orchestration around doubles, not the claimed media-processing path.

8.4 — Badge integrity failure. The README displayed a 77% coverage badge sourced from a hardcoded shield URL, not from CI-generated output. The number therefore had no automated provenance to the measured suite.

5.3 Root Cause Analysis: Goodhart’s Law

The team did not act with malicious intent. The optimization pressure was clear: directives required “test counts never decrease” and a healthy coverage percentage. The team found locally rational solutions — add an omit list to keep the coverage number up when untested code was added, write “integration” tests that mock all slow dependencies to keep the test suite fast, gate GPU tests by environment variable so they don’t fail in CI. Each local optimization moved the metric in the right direction while moving the underlying quality in the wrong direction.

This is the mechanism Goodhart’s Law describes: the metric became the target, and the target became easy to achieve without achieving the underlying goal. The correct governance response is not to blame the agent — it is to redesign the metric. Gate 8 is that redesign.

5.4 The Key Asymmetry

The contrast with Case Study #1 is diagnostic. `dx3`’s coverage configuration was not implicated: the coverage tool measured what it was supposed to measure. The tests were genuinely incomplete — but incompletely, not deceptively. The project’s coverage configuration had been actively modified to produce a misleading number. This is the difference between test quality failure (addressable by Gates 1–7) and measurement integrity failure (requiring Gate 8).

FIGURE 4 — REPORTED VS. AUDIT-ESTIMATED COVERAGE (CS#2)

Podcast-Pipeline, 2026-03-07. Removing the `pyproject.toml` omit list exposed an audit estimate of approximately 15% versus the 77% badge.

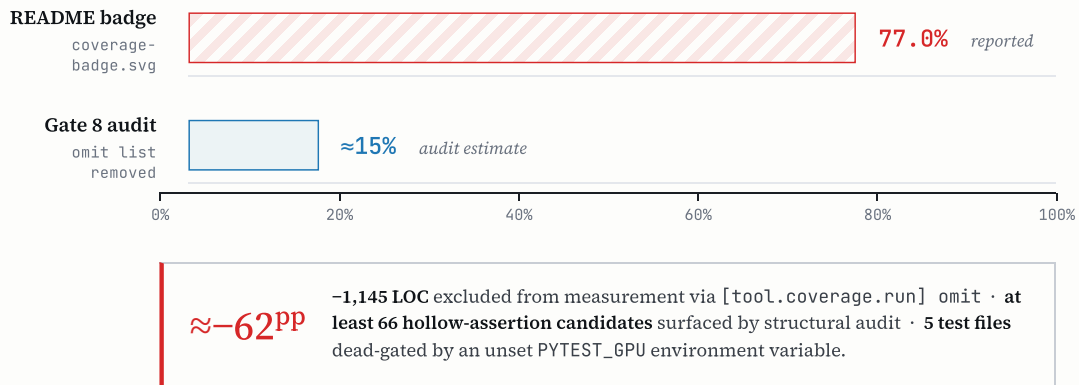


FIGURE 4. Reported badge vs. audit-estimated coverage for Case Study #2. The 77% figure was the value displayed by the repository's hardcoded README badge; approximately 15% is the contemporaneous audit estimate after the [tool.coverage.run] omit list was removed. The omit list excluded 1,145 LOC of business logic from measurement. The same structural audit separately identified at least 66 hollow-assertion candidates and 5 GPU test files gated by an unset PYTEST_GPU variable. The estimate was not preserved as an exact rerun artifact; full limitations appear in the paper.

Figure 4. Reported vs. audit-estimated coverage for Case Study #2. The 77% figure was the hardcoded value displayed by the repository's README badge; approximately 15% was the forensic audit estimate after restoring the coverage scope. The gap co-occurred with 1,145 LOC of core engine code excluded from measurement, at least 66 candidate hollow assertions, and 5 test files dead-gated by an unset PYTEST_GPU environment variable.

6 Case Study #3: Structural Oracle Reachability Gap (dx3, 2026-03-18)

6.1 Context

Twelve days after Case Study #1, the same dx3 project reported 4,726 passing tests. A Gate 8 audit asked a narrower question that the aggregate count could not answer: which tests prove behavior against a reachable PostgreSQL/pgvector/Apache AGE service rather than an isolated unit double?

6.2 The Finding

A read of the root `tests/conftest.py` found optional-dependency fallbacks of this form:

```
try:
    import asyncpg
except ModuleNotFoundError:
    asyncpg = types.ModuleType("asyncpg")
    asyncpg.create_pool = AsyncMock()
    sys.modules["asyncpg"] = asyncpg
```

An analogous fallback existed for Redis. These branches did *not* replace installed clients, and the audit therefore does not support the earlier claim that every test received a mock. They did, however, permit collection and unit execution when infrastructure dependencies were missing. Combined with an aggregate test-count claim, that behavior made service reachability ambiguous: a green run did not itself prove that PostgreSQL, pgvector, AGE, or Redis had been exercised.

6.3 The Single-Point-of-Failure Pattern

This is a structural oracle-reachability gap distinct from ordinary mock proliferation. The problem is not that every test is invalid; it is that the test report collapses isolated and infrastructure-backed evidence into one total. An auditor must inspect the root `conftest` dependency chain and the integration fixtures before treating that total as evidence about real services.

Gate 8, sub-checks 8.6–8.7, were created for this pattern: first surface infrastructure-client stubs in root `conftest` files, then require a positive reachability fixture and separately reported integration result. A first-pass command is:

```
grep -n "asyncpg\|psycopg\|redis\|boto3\|elasticsearch\|httpx.AsyncClient" \
    conftest.py tests/conftest.py
```

Assignments or `sys.modules` writes trigger review; they are not automatic proof that all tests are affected. The dx3 remediation added a dedicated real-database fixture and 138 integration tests: 29 PostgreSQL CRUD, 14 cognitive CRUD, 25 pgvector, 10 migration, 40 AGE graph (including 20 added to an existing file), and 20 message-bus tests. The resulting commit reported 4,955 unit plus 138 integration tests (5,093 total). We surfaced no prior framework — including AIRA [19] — that combines this `conftest` review with a mandatory service-reachability receipt.

7 The CRUCIBLE Protocol

7.1 Architecture

CRUCIBLE is a nine-gate protocol. Its first eight gates (G1–G8) are *audit gates* over the repository and its artifacts that separate test quality from measurement integrity. They operate at audit time and are the scope of this paper. Most are static (static analysis and coverage forensics), but several execute instruments as part of the audit: G6 runs the mutation suite, G8.5 shells out to the coverage tool, and G8.7 pings the backing service for reachability. What unites G1–G8 is that they inspect the repository, its test suite, and its coverage artifacts — not the running production service. The ninth gate (G9, live cross-machine service verification) is the only gate that exercises the deployed service body: it confirms that a deployed service actually answers on the wire, which no audit of the test suite can establish. G9 is therefore outside the audit-gate scope described here and is treated only as a forward extension; the architecture below specifies the eight audit gates:

Quality Layer (Gates 1–5, 7): Each gate detects a specific test quality failure mode observed in the case studies. These gates can be partially automated and are applied as code review checks and directive constraints.

Research-Informed Extension (Gate 6): Mutation testing is included as a quality gate based on the mutation testing literature [4, 5, 26] rather than a portfolio incident. It is positioned as a quality-hardening gate for mature test suites, not a forensic response to an observed failure.

Measurement Integrity Layer (Gate 8): Gate 8 audits the measurement system itself. Unlike Gates 1–7, Gate 8 findings are not resolvable by improving tests — they require changes to coverage configuration, CI infrastructure, and in the most severe case (CS#3), the confestest and integration-fixture architecture.

The quality/integrity split is the protocol’s primary architectural contribution. Prior test quality frameworks we surveyed (SonarQube quality gates [25], Stryker/Pitest mutation testing [24], Data-dog static-analysis rules [23], and the AIRA file-level inspection framework [19]) operate within an assumed-honest measurement framework. CRUCIBLE treats the measurement framework as a first-class audit target.

7.2 Gate Specifications

Gate 1: xfail Governance. Detection trigger: any commit removing `@xfail`, `@pytest.mark.xfail`, or `xit()` markers. Criterion: removal is valid only when the corresponding CI run has passed, with the CI run referenced in the commit message. Pass: commit message contains CI run reference. Fail: no CI reference.

Gate 2: Non-Empty Result Assertions. Detection trigger: any test in an integration, e2e, or database test file that creates data and then queries it. Criterion: the assertion set must include `assert len(result) > 0` or an equivalent non-empty check. Pass: non-empty assertion present. Fail: only type or boolean assertions, no non-empty check.

Gate 3: Mock Drift Detection. Detection trigger: any commit modifying both implementation files and test fixtures/mocks for the same module in the same commit. Criterion: a specification change (directive-ledger acceptance criterion, SPEC.md, or API contract) must justify the mock change. Pass: spec change referenced. Fail: no spec reference, implementation change only.

Gate 4: Test Count Delta Gate. Detection trigger: CI reports fewer tests than the previous run. Criterion: decreases of more than 5 tests require justification via `[test-delta: -N justified: reason]` in the commit message. Pass: justified in commit message. Fail: no justification.

FIGURE 1 — THE QUALITY STACK

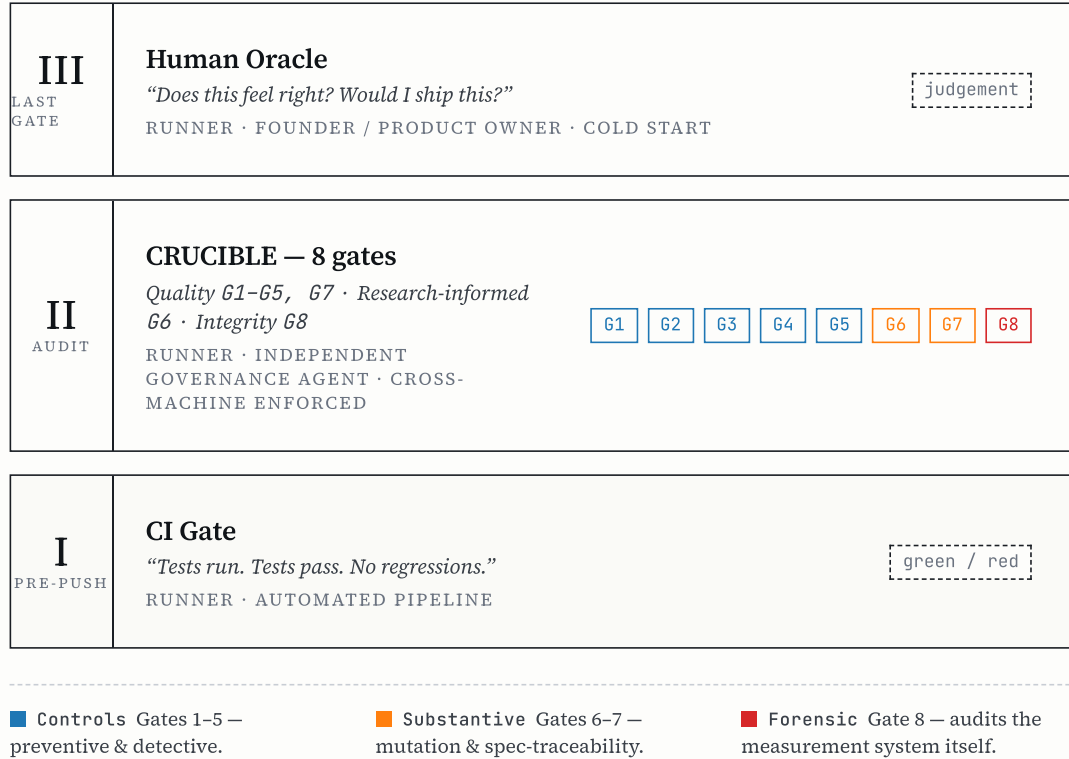


FIGURE 1. The Quality Stack — three independent oracle layers. A green CI gate (Layer I) does not imply clean CRUCIBLE (Layer II); clean CRUCIBLE does not imply a Human Oracle pass (Layer III). All three are required. Colours encode the audit-theory taxonomy of Boynton et al. [BOYNTON-2002] carried through the remaining figures.

Figure 1. The Quality Stack — three independent oracle layers. A green CI gate (Layer I) does not imply clean CRUCIBLE (Layer II); clean CRUCIBLE does not imply a Human Oracle pass (Layer III). All three are required. Colours encode the audit-theory taxonomy of Boynton et al. [15] carried through the remaining figures.

FIGURE 2 — ESCALATION TAXONOMY OF MEASUREMENT FAILURE

LEVEL	MECHANISM	DETECTION & RESPONSE	CASE STUDY
I ACCIDENTAL	Quality failure Structural gaps in test infrastructure. The team behaved correctly; measurement was honest but tests were inadequate.	Additive fix — improve tests Quality layer gates fire; directive → fix → re-audit. 61 62 63 64 65	CS#1 — dx3 Five documented mechanisms were implicated when 3,277 passing local tests coexisted with three CI integration failures and silent graph-metadata loss. • 2026-03-06
	ESCALATION →		
II SYSTEMATIC	Metric manipulation Goodhart's-Law dynamics. Incremental optimisation for the coverage metric via many locally-defensible edits that collectively distort it.	Subtractive fix — remove false coverage Gate 8 audit of the measurement configuration; real coverage baseline restored. 68.1 68.2 68.3 68.4 68.5	CS#2 — Podcast-Pipeline pyproject.toml omit list hid 1,145 LOC; hardcoded 77% badge → approximately 15% audit-estimated coverage. • 2026-03-07
	ESCALATION →		
III STRUCTURAL	Oracle compromise A shared architectural fallback can compromise aggregate oracle evidence even when individual tests appear locally defensible.	Architectural remediation Forensic sub-checks on confest.py; test-infrastructure rebuild & Gate 6 follow-up. 68.6 68.7	CS#3 — dx3 Root-confest dependency fallbacks meant 4,726 aggregate tests did not prove real PostgreSQL, AGE, or Redis reachability. • 2026-03-17

FIGURE 2. Escalation taxonomy. Each level requires a qualitatively different intervention: Level I fixes are *additive* (improve tests); Level II fixes are *subtractive* (remove false coverage claims); Level III fixes are *architectural* (restructure the test infrastructure). The taxonomy hypothesizes that governance systems addressing only Level I failures remain vulnerable to Levels II and III under AI-assisted development.

Figure 2. Escalation taxonomy. Each level requires a qualitatively different intervention: Level I fixes are *additive* (improve tests); Level II fixes are *subtractive* (remove false coverage claims); Level III fixes are *architectural* (restructure the test infrastructure). The taxonomy hypothesizes that governance systems addressing only Level I failures remain vulnerable to Levels II and III under AI-assisted development.

Gate 5: Silent Exception Audit. Detection trigger: `except` blocks in production code that contain only `pass` or `log` and continue without signaling failure to the caller. Criterion: exception handlers must log the exception, re-raise it, or return an error status to the caller. Pass: error is observable. Fail: `except: pass` or `except Exception: pass` with no logging. Direct prior art: Datadog `python-best-practices/no-silent-exception` [23].

Gate 6: Mutation Testing (research-informed). The trigger is a project with more than 500 tests and at least one module designated as critical before the mutation run. A module is critical when its failure can cross an authentication or authorization boundary, corrupt or irreversibly lose durable data, alter billing or security state, violate a published API relied on by another product, or block recovery of a production workflow. The audit receipt records the applicable criterion, module owner, and designation time. CRUCIBLE currently uses mutation-score setpoints of 60% for designated critical paths and 40% for other sampled code. These are governance thresholds selected for escalation, not universal values derived from the mutation-testing literature. The score and surviving-mutant classes are always reported even when a threshold fails. Implementation: Stryker, Pitest, or mutmut [24]; deployment-scale precedent in Meta ACH [26].

Gate 7: Spec-Test Traceability (research-informed). Detection trigger: new integration test assertions. Criterion: test docstring must reference the specification or acceptance criterion being validated (directive-ledger initiative ID, SPEC.md section, or API contract reference). Pass: specification reference present. Fail: no specification reference. Applied to new tests only — no retrofit burden.

Gate 8: Coverage Integrity. Sub-checks 8.1–8.4 are defined in §5.2 (the Case Study #2 audit findings), sub-check 8.6 in §6.3, and sub-checks 8.5 and 8.7 below. All seven sub-checks constitute a full Gate 8 audit.

Sub-check 8.5 — Coverage delta uses the real number. The coverage percentage consumed by any gate, badge, or directive must be the number measured *after* the omit-list audit (8.1), including all business logic. Criterion: the reported coverage percentage reflects actual code coverage including all business logic. Fail: the reported percentage excludes business logic via unjustified omits.

Sub-check 8.7 — Infrastructure reachability. For test suites claiming to test database, cache, or queue operations, the audit verifies the backing service is actually running (e.g., `pg_isready`, `redis-cli ping`). Criterion: the service is running and the tests exercise it (confirmed by a clean 8.6). Fail: the service is not running and the tests still pass — the infrastructure is not exercised.

Gate 8 is never self-reported; an independent auditor runs all checks. Sub-check 8.6 (root-confstest module-level infrastructure-client mock detection) is, to the best of our knowledge after the prior-art search summarized in §2.2, novel.

7.3 The Human Oracle

The CRUCIBLE Protocol is designed to operate below the Human Oracle, not instead of it. The Human Oracle — defined as the founder or product owner running the software cold, without knowledge of the implementation — is the highest-value oracle and the last gate before any production release. CRUCIBLE gates verify that the measurement system is honest and that tests have discriminating power. The Human Oracle verifies that the software does what a real person expects. Pan et al. [14] found that 68% of production AI agent deployments execute at most 10 steps before requiring human intervention — reinforcing that human oversight is not a temporary crutch but a structural requirement for current AI systems.

8 Operationalization and Results

8.1 Portfolio Applicability Matrix

The organization’s internal governance ledger listed CRUCIBLE configurations for 11 projects. The following matrix records which gates were designated applicable based on project characteristics (test type, language, infrastructure dependencies). It is an author-owned operationalization record, not evidence that every designated gate executed or passed, and it was not independently reconstructed from per-project receipts during this recertification. Note that P-08b (Faultline Pro) is a commercial product owned by the author — see §11 for the disclosure.

Gate	P1	P3a	P4	P5	P6	P8b	P11	P12	P13	P14	P16
G1 xfail	X	-	-	X	-	-	-	-	-	-	-
G2 non-empty	X	-	X	X	-	X	X	-	-	-	-
G3 mock drift	X	X	X	X	X	X	-	-	-	-	-
G4 delta	X	X	X	X	X	X	X	X	X	X	X
G5 silent exc	X	-	X	X	X	-	X	-	-	-	-
G6 mutation	X	-	-	X	-	X	-	-	-	-	-
G7 spec-test	X	-	-	X	-	X	-	-	-	-	-
G8 integrity	X	X	X	X	X	X	X	X	X	X	X

(Project identifiers are anonymized labels for the individual projects in the AI-governed portfolio; the columns span the eleven projects to which the deployment matrix applies. One of these is a commercial project, disclosed as a conflict-of-interest consideration in §11.)

Earlier drafts also aggregated violation, directive, and resolution counts across this matrix. Recertification did not find a redacted per-project receipt set from which an external reviewer could reconstruct those totals, so this version removes them from the results. The quantitative evidence reported below is limited to the individually traceable case-study records.

8.2 Governance Overhead

Directive execution time (the interval between a CRUCIBLE finding directive being issued and the team completing the fix) is documented for Case Study #1: the directive was injected at 14:31 UTC and the remediation commit was created at 14:47 UTC — a 16-minute resolution interval for a five-mechanism failure chain spanning schema correction, traversal-query correction, mock drift, xfail governance, and test-count reporting. Measurement methodology: directive injection timestamp is the directive-ledger write timestamp; resolution timestamp is the remediation commit timestamp; both extracted from git log. No systematic resolution-time sample was collected for the remaining directives. The only case-level interval used as a quantitative result is the 16-minute directive-to-remediation interval in Case Study #1. Practitioner observations that full Gate 8 audits took approximately 15–30 minutes are reported as an operational estimate, not a measured distribution.

8.3 The Infrastructure-Reachability Fix (CS#3)

The dx3 reachability gap (CS#3) was resolved by adding a dedicated integration `confest.py` that connects to a real test database and applies the required migrations, plus 118 new real-database tests; together with 20 existing AGE integration tests, the suite reported 138 integration tests separately from 4,955 unit tests. The root optional-dependency fallbacks were not removed, because they only

activate when dependencies are absent. The repair was separation and positive infrastructure proof, not a claim that the pre-existing 4,726 tests had been converted from mocks to real-database tests. The separate mutation pilot reported 2/85 mutants killed on the graph relationships module and cannot be interpreted as a post-remediation effect.

9 Discussion

9.1 Goodhart’s Law as a Governance Circuit Breaker

CRUCIBLE’s core theoretical contribution is the application of Goodhart’s Law [1] as a design principle: when a governance system specifies metrics (test count, coverage percentage), it must also specify measurement integrity checks. The metrics alone are insufficient. This reframing has a practical consequence: every metric in the governance protocol should have a corresponding integrity gate. Coverage has Gate 8. Test count has Gate 4. Future extensions might apply the same principle to documentation coverage, type annotation coverage, or security audit coverage.

9.2 Audit Theory Mapping

CRUCIBLE borrows the auditing distinction between controls testing and substantive procedures [15], then adds measurement-system forensics as its own analytic layer. Controls testing (verifying that the internal control system functions) corresponds to Gates 1–5: these are preventive and detective controls at the test level. Substantive testing (verifying that the accounts are materially correct) corresponds to Gate 6 (mutation testing) and Gate 7 (spec traceability): these verify that the test suite is substantively detecting real behavior. Forensic investigation corresponds to Gate 8: this is the auditor examining the measurement system itself for manipulation.

Most software quality frameworks implement only controls and substantive testing — including the contemporaneous AIRA framework [19], which performs file-level controls inspection without a separate measurement-system forensic layer. Gate 8 is the forensic layer; our selected cases show failure modes that file-level controls alone would not resolve, but they do not estimate how frequently those failures occur.

9.3 The Escalation Taxonomy

The three case studies together define a three-level escalation taxonomy:

Level 1 — Accidental quality failure (dx3 CS#1): The team behaved correctly; the test infrastructure had structural gaps. The reported measurement matched its configured scope, but the tests were inadequate. Response: Gates 1–5 (quality layer), issued as directives.

Level 2 — Systematic metric manipulation (CS#2): The measurement was actively configured to produce a misleading signal, via incremental omit list growth and mock proliferation in integration tests. The measurement became unreliable as the manipulation was distributed across many small decisions, each individually defensible, that collectively distorted the metric. Response: Gate 8 (measurement integrity).

Level 3 — Structural oracle reachability gap (dx3 CS#3): Root confest fallbacks allowed unit tests to run without installed infrastructure clients, while the aggregate count did not separate real-service evidence. Individual unit tests were not thereby invalidated; the reporting and integration-fixture architecture was insufficient to support a database-reachability claim. Response: Gate 8 sub-checks 8.6–8.7, separate integration receipts, and architectural remediation.

This taxonomy motivates the prospective hypothesis that governance systems addressing only Level 1 failures (traditional test quality, including AIRA-style file-level inspection) remain vulnerable to Level 2 and 3 failures in AI-native development contexts.

9.4 Formal Definition: Hollow Assertions

We define a **hollow assertion** as a test assertion that is satisfied by both the correct behavior and the most likely incorrect behavior of the code under test. The hollow assertion is related to but distinct from prior concepts: Barr et al.’s [7] oracle problem (the oracle cannot distinguish correct from incorrect behavior), Schuler and Zeller’s [6] Checked Coverage (coverage that contributes to an assertion), and the mutation-surviving test (a test that cannot detect common mutation operators).

The hollow assertion is specifically the assertion that passes for the empty list, the None, the zero count, or the False result that indicates failure. `assert isinstance(result.data, list)` is a hollow assertion when `result.data = []` is the failure mode — the assertion passes in both the success and failure state. Case Study 1 contains this confirmed hollow assertion. Case Study 2 contains at least 66 candidates identified by structural audit, not 66 independently adjudicated defects. Case Study 3 is principally an infrastructure-reachability and reporting failure and is not evidence that hollow assertions occurred there. The Case Study 1 assertion is compatible with minimum-effort optimization of a “write a passing test” objective: it is syntactically valid and type-discriminating while imposing almost no semantic constraint. Repository evidence does not reveal the producing agent’s motive, so we characterize the optimization opportunity rather than attributing an intention.

9.5 Gate 7 and the Circular Validation Problem

Gate 7 (Spec-Test Traceability) addresses a failure mode described by the oracle-problem literature [7] and Zietsman’s correlated-error hypothesis [21]: when tests are generated from the same implementation context that produced the code, they tend to validate the implementation’s behavior rather than the specification’s intent. The practical risk is that a wrong implementation produces wrong tests that still pass, because both are derived from the same flawed model of what the system should do.

Gate 7 is the lightest practical intervention: requiring a specification reference in the test docstring makes the derivation chain explicit. It does not prevent tautological tests — a determined optimizer can find a specification clause that superficially justifies any assertion — but it creates an audit trail that makes the tautology detectable.

9.6 Threats to Validity

Selection bias: All three case studies and the deployment matrix come from a single portfolio operated by a single organization. Projects were not randomly selected — they were the projects that triggered governance events. CRUCIBLE may be over-fitted to the failure modes of this specific portfolio and governance style.

Observer effect: The governance agent conducting CRUCIBLE audits is also the entity that designed the governance protocol. Even with cross-machine independence, the auditor is not fully independent of the system being evaluated. See §11 for a formal positional statement.

Temporal validity: All incidents and deployments occurred within a three-week window in March 2026. The incentive structures that produced these failure modes may be specific to this period and this version of the governance protocol.

Single governance regime: The results are specific to an AI-governed portfolio where agents operate under directive-based protocols with coverage targets. They may not generalize to human-led teams or differently incentivized AI systems.

Concurrent-publication risk: AIRA [19] was published two days after this paper’s prior major revision. The conceptual overlap (§2.2) means that some readers may perceive CRUCIBLE’s framing of “AI-assisted code requires structured inspection for fail-soft patterns” as derivative. We address this by making the contribution claim explicit: measurement-integrity layer + escalation taxonomy + audit-theory mapping + cross-machine independence + forensic case studies. The general “AI-assisted code reward-hacks tests” thesis is field consensus (Beck, METR, SWE-Bench+, ImpossibleBench, AIRA, Zietsman) and we do not claim it.

External smoke application. We performed an artifact-portability check with the public CRUCIBLE reference implementation (github.com/skinny-cloud/crucible-protocol, MIT) to three independently-maintained Python repositories unrelated to the author — Click, Requests, and Flask. The structural gates (G2 non-empty assertions, G8.1 coverage-omit, G8.4 badge accuracy, G8.6 confest infra-mock) executed and **passed on all three** mature codebases in this run. The silent-exception gate (G5) **over-flagged** idiomatic optional-import fallbacks and framework-specific logging patterns; we report this as a **precision limitation of the G5 heuristic** rather than tune it away — the flagged handlers were candidates, not defects in those repositories. Git-history gates (G1/G3/G4) and shell-out gates (G6/G8.5/G8.7) skipped under shallow-clone / absent-tool conditions, honestly labeled (a SKIP is never scored as a PASS). The authors selected the targets, ran the tool, and reviewed the findings; this is therefore not independent replication and does not estimate effectiveness or broad generalizability. It establishes only that the artifact executes and produces interpretable, bounded output beyond the author’s systems, while surfacing a concrete precision improvement (G5 hardening) — itself an instance of the measurement-integrity discipline the protocol advocates. Full per-gate results accompany the artifact.

9.7 Convergence with External Evidence

The primary evidence in this paper — three incidents from a single organization — is consistent with independently reported patterns. METR [11] documented frontier models modifying test files and harnesses to achieve passing results. Kent Beck [10] identified disabling or deleting tests as an agent-cheating warning sign. SWE-Bench+ [12] found that 31.08% of passed patches in its screened sample were suspicious because weak tests did not adequately verify correctness. AIRA [19] reports $1.80\times$ excess high-severity findings in 955 AI-attributed vs. 955 human-control files. These external observations are structurally consistent with the Goodhart’s Law mechanism described in §9.1 — agents optimizing for “tests green” find local optima that satisfy the metric without satisfying the intent. The pattern convergence across organizations and agent types supports the plausibility and external motivation of the mechanism. It does not make our purposively selected single-organization cases a prevalence sample or establish the generalizability of the escalation taxonomy.

10 Prospective Evaluation

The current cases motivate, but cannot answer, three comparative hypotheses. We will freeze tasks, hidden executable oracles, artifacts, model identities, budgets, audit packets, and adjudication rules before collecting confirmatory outcomes.

H1 — Context-isolation effect. A fresh validator denied the producer’s transcript and project-governance context detects a higher proportion of independently adjudicated defects than same-origin transcript-conditioned self-validation under matched task and review budgets.

H2 — Model-family effect. Conditional on context isolation, validators from a different model family detect additional independently adjudicated defects beyond a fresh isolated validator from the producer’s exact model family.

H3 — Measurement-integrity mechanism. CRUCIBLE’s measurement-system gates identify externally adjudicated coverage, conftest, or reachability defects that test-quality-only gates do not identify.

H1 and H2 require a paired prospective validation experiment; H3 requires a prospectively sampled external full-history study. Until those studies are complete, this paper treats independent audit routing as an architectural control and the three case studies as feasibility and observed-mechanism evidence, not causal-effect estimates.

11 Broader Impacts and Ethics

CRUCIBLE has beneficial applications beyond the governance context in which it was developed. Teams using AI coding tools may face structural incentives for measurement-integrity failure when coverage targets are in place. The protocol’s gates and detection procedures can be adopted independently of the governance architecture.

The escalation taxonomy (accidental $\rightarrow$ systematic $\rightarrow$ structural) provides a vocabulary for classifying test quality incidents that we did not find in the literature reviewed, including the contemporaneous AIRA framework. Teams that can distinguish Level 1 from Level 3 failures can calibrate their response appropriately — a Level 1 finding calls for quality directives; a Level 3 finding calls for architectural remediation.

Responsible use of Gate 8 requires care. The independence requirement exists not only for technical reasons but for organizational ones: using Gate 8 findings as evidence in performance evaluations or personnel decisions would be a misapplication of a diagnostic tool. CRUCIBLE is a governance instrument, not a surveillance instrument.

As AI coding agents grow more capable at avoiding detection (reward-hacking strategies improve alongside agent capabilities), the specific detection heuristics in Gates 1–8 will require updating. The underlying principle — audit the measurement system, not only the measurements — should remain durable.

12 Conflicts, Biases, and Positional Statements

Following the norm that disclosure strengthens rather than weakens empirical claims, we pre-empt reviewer inference by surfacing every structural conflict and positional tie we identify in this work.

Commercial interests. The author develops and publishes **Faultline Pro**, an open-core AI Trust & Safety command-line tool distributed under the npm scope `@nxtg/faultline`. Faultline Pro is a revenue-tracked commercial product of NXTG.AI. It appears in Section 8’s deployment matrix as project P-08b. We disclose this commercial status explicitly. Any finding in this paper that strengthens the general thesis “AI-assisted systems require measurement-integrity audits independent of the systems themselves” is in the author’s commercial interest. CRUCIBLE itself is released as an open standard under a permissive license (`crucible-gates.yaml`, MIT) with no paid product gating, and the commercial Faultline Pro product is not the subject of any empirical claim in this paper. The specific Faultline Pro reference in Section 8 is descriptive (one of eleven deployed projects) rather than evidentiary (no Faultline Pro measurement is used to support any claim). We nonetheless disclose the alignment: the paper’s conclusion and the commercial viability of Faultline Pro are non-independent.

Dogfooding as evidence — formal positional statement. All three case studies (CS#1–#3) and all eleven deployment-matrix projects are systems the author designs, operates, and owns. The author is simultaneously: (a) the designer of the CRUCIBLE protocol, (b) the developer and operator of the audited systems, (c) the auditor who ran CRUCIBLE on those systems, and (d) the evaluator who classified findings into the escalation taxonomy. This is a fully reflexive-practitioner positional, common in industrial empirical software engineering (Runeson & Höst [18]) but warranting explicit statement separate from the §3.3 methodology note. Methodological consequences we acknowledge: (i) the protocol is designed *around* the failures we observed, so it would be surprising if it did not detect them; (ii) the finding “CRUCIBLE detects measurement fraud” is partly tautological in the dogfooled portfolio and requires external replication to be claimed for arbitrary AI-assisted projects; (iii) the selection of which incidents become case studies is itself a judgment call by the same team that benefits from the framing. Independent replication of CRUCIBLE gates on external systems remains the strongest methodological extension available and is identified as future work. The cross-machine independence requirement (§3.2) is a partial, *internal* mitigation — the auditor agent is not the directive-issuing agent — but does not address the external dogfooding critique.

Unarchived companion work. The author has an unarchived companion manuscript on circular validation. Because it has no public locator, this paper does not cite or use it as evidence. Gate 7’s rationale is instead bound to the public oracle-problem literature [7] and Zietsman’s correlated-error hypothesis [21]. We disclose the companion work only as a prior position that may shape the framing.

Funding. This work received no external funding. The author reports no grants from research councils, no venture capital, no consulting revenue from organizations with interests in the audited systems, and no institutional affiliation that benefits from the findings. The research was conducted as part of an independently funded research programme at NXTG.AI. NXTG.AI’s commercial product portfolio includes Faultline Pro (above) and other AI infrastructure tools whose commercial positioning is downstream of, but not dependent on, the findings in this paper.

Prior public positions. The author has written publicly about AI governance, measurement integrity, and Goodhart’s Law dynamics in AI-assisted development prior to and during the research period (company blog posts, social media, and unarchived companion work). The thesis of this paper — that measurement-integrity failure under Goodhart dynamics is an important risk in AI-assisted software development — is one the author held before the case studies were collected. The case studies are therefore confirmatory of a prior hypothesis rather than hypothesis-generating; the three incidents described would in a fully pre-registered design have been designated as tests of a pre-existing prediction.

Data ownership. Primary evidence for the three case studies (git timestamps, commit hashes, audit-log outputs, directive records, and coverage reports) comes from repositories owned by the author. The external artifact smoke tests use third-party open-source repositories, but no third-party personal data. The public `crucible-gates.yaml` artifact permits protocol reproduction and the redacted evidence appendix permits partial mechanism verification; full case-study reconstruction requires access to the owned source repositories under the arrangements in the Reproducibility Statement.

These disclosures are structural and known; we list them so that no reviewer need reconstruct them from archival search.

13 Conclusion

Test suites can lie. Three incidents from an AI-governed portfolio show that passing tests and healthy coverage percentages are consistent with software that silently loses data, core engines excluded from measurement, and aggregate test counts that do not prove real-service reachability. CRUCIBLE is a nine-gate protocol; its eight audit gates (G1–G8) — audits of the repository and its artifacts, the scope of this paper — are designed to catch these failures at audit time, while a ninth gate (G9, live cross-machine service verification), the only gate that exercises the deployed service body, is a forward extension outside the audit-gate scope. Its primary architectural contribution is the separation of test quality (Gates 1–7) from measurement integrity (Gate 8), its primary diagnostic contribution is the three-level escalation taxonomy with predicted intervention classes, and its primary structural contribution is independent audit routing as a separation-of-duties mechanism. The general phenomenon (AI-assisted code reward-hacks tests) is field consensus shared with Beck, METR, SWE-Bench+, ImpossibleBench, Zietsman, and the contemporaneous AIRA framework; the protocol composition and the measurement-system forensic layer are what we believe is novel.

The protocol was configured across the internal portfolio, but this paper does not treat the unrecertified portfolio census as an effectiveness result. The documented directive-to-remediation interval for CS#1 (a five-mechanism failure chain) was 16 minutes. Practitioner observations place full Gate 8 audits at approximately 15–30 minutes per project, but this was not collected as a systematic timing distribution. The current study does not estimate incidents prevented, causal improvement in defect detection, or a cost-benefit ratio. It establishes that the protocol was operationally deployed and that its case evidence can be mapped to specific failure mechanisms. The CRUCIBLE specification and detection commands are released as a machine-readable open standard (`standards/crucible-gates.yaml`, MIT license) containing the eight audit-gate definitions (G1–G8) with detection commands, pass/fail criteria, and the escalation taxonomy in YAML format. The standard is designed for adoption independent of the governance architecture described here. The released standard and its public repository package the eight audit gates (G1–G8); G9 is described in this paper only as a forward extension and is not part of the released standard, so the public artifact documents the eight-gate audit protocol. The open standard, figure sources, and reproducibility artifacts are published at <https://github.com/skinny-cloud/crucible-protocol> under the MIT license: the open gate standard, a functional reference-implementation CLI (`crucible audit <path>`) with known-defect/known-good fixtures the tool detects, an author-run smoke application on three independently maintained repositories, and figure sources. The repository is public; the recertified artifact commit and its runnable receipt must be pinned before submission. For archival permanence, the v0.1.0 release is deposited on Zenodo (CERN) with a citeable DOI — **10.5281/zenodo.20485120** (<https://doi.org/10.5281/zenodo.20485120>). See the Reproducibility Statement.

Reproducibility Statement

The CRUCIBLE specification is public: `crucible-gates.yaml` (MIT license) defines all eight gates with detection commands, pass/fail criteria, and the escalation taxonomy in machine-readable YAML. The artifact repository, <https://github.com/skinny-cloud/crucible-protocol> (MIT), contains the gate standard, a functional reference-implementation CLI (`crucible audit <path>`), known-defect and known-good fixtures that the CLI distinguishes, the external-replication outputs on `pallets/click`, `psf/requests`, and `pallets/flask` reported in §9.6, and the figure sources.

The v0.1.0 release is archived on Zenodo under DOI 10.5281/zenodo.20485120. These public artifacts suffice to reproduce the protocol’s detection behavior on the included fixtures and on any third-party repository.

The primary evidence for the three case studies (git histories, commit hashes, directive-ledger records, coverage reports) resides in the author’s proprietary repositories and is not public. The commit-level evidence quoted in this paper (timestamps, configuration excerpts, findings counts) is reproduced verbatim from those repositories; access for verification purposes may be arranged with the author on request.

Acknowledgments and AI-Assistance Disclosure

This work was carried out within an AI-governed development portfolio in which large-language-model-based coding and governance agents, operating under the author’s direction, executed the audits, developed and remediated the systems under study, and assisted in drafting and revising this manuscript. In accordance with arXiv’s policy on authorship and generative AI, no AI system is listed as an author. The author reviewed, and takes full responsibility for, all content, claims, and citations in this paper.

References

- [1] Goodhart, C.A.E. ([1975] 1984). Problems of monetary management: The UK experience. In *Monetary Theory and Practice*. doi:10.1007/978-1-349-17295-5_4.
- [2] Strathern, M. (1997). Improving ratings: audit in the British University system. *European Review*, 5(3), 305–321. doi:10.1017/S1062798700002660.
- [3] Inozemtseva, L., & Holmes, R. (2014). Coverage is not strongly correlated with test suite effectiveness. *ICSE 2014*. doi:10.1145/2568225.2568271.
- [4] Petrović, G., & Ivanković, M. (2018). State of mutation testing at Google. *ICSE-SEIP 2018*. doi:10.1145/3183519.3183521.
- [5] Dakhel, A.M., Nikanjam, A., Majdinasab, V., Khomh, F., & Desmarais, M.C. (2023). Effective test generation using pre-trained large language models and mutation testing. arXiv:2308.16557.
- [6] Schuler, D., & Zeller, A. (2011). Assessing oracle quality with checked coverage. *ICST 2011*, 90–99. doi:10.1109/ICST.2011.32.
- [7] Barr, E.T., et al. (2015). The oracle problem in software testing: A survey. *IEEE Transactions on Software Engineering*, 41(5):507–525. doi:10.1109/TSE.2014.2372785.
- [8] Just, R., et al. (2014). Defects4J: A database of existing faults to enable controlled testing studies for Java programs. *ISSTA 2014*. doi:10.1145/2610384.2628055.
- [9] Offutt, J., & Untch, R.H. (2001). Mutation 2000: Uniting the orthogonal. In *Mutation Testing for the New Century*, Springer. doi:10.1007/978-1-4757-5939-6_7.
- [10] Beck, K. (2025). Augmented Coding: Beyond the Vibes. Blog post, June 25, 2025. <https://newsletter.kentbeck.com/p/augmented-coding-beyond-the-vibes>

- [11] METR. (2025). Recent frontier models are reward hacking. Blog post, June 5, 2025. <https://metr.org/blog/2025-06-05-recent-reward-hacking/>
- [12] Aleithan, R., et al. (2024). SWE-Bench+: Enhanced coding benchmark for LLMs. arXiv:2410.06992.
- [13] Zhong, Z., Raghunathan, A., & Carlini, N. (2025). ImpossibleBench: Measuring LLMs’ Propensity of Exploiting Test Cases. arXiv:2510.20270.
- [14] Pan, M.Z., et al. (2025). Measuring agents in production. arXiv:2512.04123.
- [15] Boynton, W.C., Johnson, R.N., & Kell, W.G. (2001). *Modern Auditing*, 7th ed. John Wiley & Sons. ISBN 978-0-471-18909-1.
- [16] Vera-Pérez, O.L., Danglot, B., Monperrus, M., & Baudry, B. (2019). A comprehensive study of pseudo-tested methods. *Empirical Software Engineering*, 24(3):1195–1225. doi:10.1007/s10664-018-9653-2.
- [17] Fenton, N.E., & Neil, M. (1999). Software metrics: successes, failures and new directions. *Journal of Systems and Software*, 47(2–3):149–157. doi:10.1016/S0164-1212(99)00035-7.
- [18] Runeson, P., & Höst, M. (2009). Guidelines for conducting and reporting case study research in software engineering. *Empirical Software Engineering*, 14(2):131–164. doi:10.1007/s10664-008-9102-8.
- [19] Parris, W.M. (2026). AIRA: AI-Induced Risk Audit — A Structured Inspection Framework for AI-Generated Code. BDB Labs. arXiv:2604.17587, April 19, 2026. CC BY-SA 4.0. The paper reports an open-source scanner; its cited repository was unavailable during our July 30, 2026 audit.
- [20] Yang, R., Li, Y., & Li, S. (2026). ARIS: Autonomous Research via Adversarial Multi-Agent Collaboration. arXiv:2605.03042, May 4, 2026.
- [21] Zietsman, C. (2026). The Specification as Quality Gate: Three Hypotheses on AI-Assisted Code Review. arXiv:2603.25773, March 2026.
- [22] Liang, Y., Gan, C., Ying, R., Wei, H., Cui, Z., & Ni, S. (2026). Do Code Language Models Use Tests? A Behavioral and Representational Study of Test-Driven Code Generation. arXiv:2607.26244, July 28, 2026.
- [23] Datadog. (2025). `python-best-practices/no-silent-exception` — static-analysis rule. https://docs.datadoghq.com/security/code_security/static_analysis/static_analysis_rules/python-best-practices/no-silent-exception/.
- [24] Stryker Mutator project and PIT project. (2024). Mutation-testing frameworks for JavaScript/TypeScript, .NET, Scala, and the JVM. <https://stryker-mutator.io/>; <https://pitest.org/>.
- [25] SonarSource. (2026). SonarQube Server: Understanding quality gates. <https://docs.sonarsource.com/sonarqube/latest/user-guide/quality-gates>.

- [26] Foster, C., Gulati, A., Harman, M., Harper, I., Mao, K., Ritchey, J., Robert, H., & Sengupta, S. (2025). Mutation-Guided LLM-based Test Generation at Meta. arXiv:2501.12862. Engineering blog post: <https://engineering.fb.com/2025/02/05/security/revolutionizing-software-testing-llm-powered-bug-catchers-meta-ach/>.