
SENTINEL: A Mid-Reasoning Interception Framework for Auditing Medical AI Agents in Production

Andrew Espira*

Department of Data Science, Saint Peter's University
aespira@saintpeters.edu

Abstract

As medical AI agents accelerate administrative and clinical workflows, their speed advantage introduces a structural risk: agents can reason incorrectly faster than humans can intervene. Existing AI safety approaches operate either post-hoc—auditing decisions after they are made—or at training time, leaving a critical gap in runtime production deployments. We present SENTINEL, a mid-reasoning interception framework that audits agent reasoning traces before irreversible healthcare decisions are executed. SENTINEL evaluates evidence quality across four dimensions—completeness, currency, weighting, and contradiction—producing a composite confidence score that gates agent actions through three intervention tiers: auto-proceed (≥ 0.85), notify-proceed (0.60–0.85), and escalate (< 0.60). We integrate SENTINEL into a healthcare revenue cycle management (RCM) agent platform built on the Model Context Protocol (MCP) across five data sources: CMS Coverage, CMS Fee Schedule, Medical Coding, NCCI Edits, and Denial Codes. In a controlled benchmark evaluation across 950 agent reasoning traces spanning three agent types—CodeResearchAgent, MedicalNecessityAgent, and DenialResolutionAgent—SENTINEL achieves a 55/18/27% green/yellow/red tier distribution, P99 audit latency of 15,006 ms, and identifies 9.3% of traces with evidence quality issues that the agents' own confidence scores did not flag. Critically, SENTINEL detected 249 security-relevant reasoning events—including PHI exposure in agent reasoning, unsafe failover suggestions, and unsupervised treatment-affecting decisions—that no existing agent self-confidence mechanism surfaces. These results demonstrate that runtime reasoning interception is a necessary and practical complement to model-level confidence scoring in production medical AI deployments.

1 Introduction

Healthcare billing generates \$262 billion in annual administrative waste CAQH [2022]. Claim denial rates have reached 11.8%, costing an average of \$57 per rework CAQH [2023]. AI agent systems promise to reduce the 30–60 minutes of manual research required per billing error to under 60 seconds by orchestrating structured access to CMS coverage databases, fee schedules, NCCI edits, and payer policies through standardized interfaces such as the Model Context Protocol (MCP) Anthropic [2024].

This speed advantage, however, creates a structural safety problem. When an agent reasons across four data sources and produces a claim submission recommendation in 60 seconds, there is no meaningful opportunity for human review before the action is taken. Claim submissions are largely

*ORCID: 0009-0002-9196-8094. Correspondence: aespira@saintpeters.edu

irreversible: once submitted, incorrect claims generate denial codes, trigger payer flags, and require expensive manual rework. The cost of an incorrect agent decision is not caught at inference time—it is caught weeks later when the explanation of benefits arrives.

1.1 The Gap in Existing Approaches

Current AI safety approaches fail to address this operational gap in three ways:

Post-hoc auditing examines decisions after they are executed Ribeiro et al. [2016], Lundberg and Lee [2017]. For reversible decisions, this is acceptable. For healthcare billing, it is not—the damage is already done when the audit runs.

Training-time alignment Ouyang et al. [2022], Bai et al. [2022] improves model behavior in aggregate but cannot account for runtime context: stale data sources, incomplete evidence sets, or domain-specific contradictions that emerge only during a specific claim workflow.

Agent self-confidence scores are present in most agent frameworks Anthropic [2025] but measure the model’s internal certainty, not the quality of the evidence the model is reasoning over. An agent can be highly confident about an incorrect conclusion when its evidence is stale or incomplete.

1.2 Contributions

We present SENTINEL, a runtime reasoning interception framework designed for production medical AI deployments. Our contributions are:

1. **Mid-reasoning interception architecture:** A hook-based system that intercepts agent reasoning traces before irreversible actions execute, requiring no modification to the underlying agent or model.
2. **Four-dimension evidence quality rubric:** A structured evaluation framework assessing completeness, currency, weighting, and contradiction of agent reasoning traces, operationalized for healthcare RCM workflows.
3. **Tiered confidence gating:** Three intervention tiers (auto-proceed, notify-proceed, escalate) with empirically derived thresholds, integrated into `agentctl`—an open-source `kubectl`-style CLI for managing AI agent deployments, providing state management, observability, and confidence gating.
4. **MCP-native integration:** SENTINEL deploys as an interceptor layer over MCP tool calls, requiring no changes to existing MCP server implementations.
5. **Benchmark evaluation:** Assessment on 950 healthcare billing agent reasoning traces across three agent types on an MCP-based RCM agent platform, with full observability via OpenTelemetry and Prometheus.
6. **Security-aware reasoning auditing:** Detection of PHI exposure, unsafe failover logic, and treatment-affecting decisions within agent reasoning traces—categories invisible to traditional output monitoring.

2 Background and Related Work

2.1 Medical AI Agents and Safety

The deployment of AI agents in clinical and administrative healthcare settings has accelerated with the availability of capable large language models Singhal et al. [2023], Nori et al. [2023]. Recent work has demonstrated strong performance on medical coding Ji et al. [2024], prior authorization Pandey et al. [2024], and claims adjudication Hendricks-Sturup et al. [2024]. Safety concerns in medical AI have focused primarily on diagnostic accuracy and clinical decision-making Obermeyer et al. [2019], with limited attention to the administrative billing pipeline where errors are costly but not life-threatening.

2.2 Runtime Monitoring for ML Systems

Sculley et al. [2015] identified the absence of prediction quality monitoring as hidden technical debt in ML systems. Subsequent work on data drift detection Rabanser et al. [2019] and operational ML Paleyes et al. [2022] has established monitoring as a first-class concern, but focuses on model outputs rather than reasoning quality. SENTINEL extends this reliability-first philosophy to agent reasoning quality, treating evidence trustworthiness as a monitorable signal analogous to data drift or prediction quality.

2.3 Model Context Protocol

The Model Context Protocol Anthropic [2024] provides a standardized interface for AI agents to access external data sources through typed tool calls. MCP’s structured tool invocation pattern creates natural interception points: every tool call has defined inputs and outputs, enabling SENTINEL to audit the reasoning that precedes a call without modifying the underlying server.

2.4 Chain-of-Thought and Reasoning Transparency

Chain-of-thought prompting Wei et al. [2022] and its variants produce reasoning traces that are auditable by design. Recent work on process-level reward models Lightman et al. [2023] and reasoning verification Uesato et al. [2022] has explored evaluating reasoning quality, but in offline/training settings. SENTINEL applies reasoning quality evaluation at runtime, in the production agent inference path.

3 The SENTINEL Framework

3.1 Problem Formulation

Let $\mathcal{A}$ be a medical AI agent executing a workflow W consisting of a sequence of tool calls $t_1, t_2, \dots, t_n$ against MCP servers $\mathcal{M}$. Each tool call t_i is preceded by a reasoning trace r_i in which the agent plans its action. Let t_k be a *high-stakes* tool call—one whose execution is difficult or impossible to reverse (e.g., `submit_claim`).

The goal of SENTINEL is to evaluate the quality of reasoning trace r_k before t_k executes, and gate execution based on a composite evidence quality score $s \in [0, 1]$.

3.2 Evidence Quality Dimensions

SENTINEL evaluates each reasoning trace across four dimensions:

Completeness (c_1): Are all required evidence fields present before the decision? In healthcare billing, this includes patient eligibility, diagnosis codes, procedure codes, and payer-specific requirements. Formally, $c_1 = |E_{\text{present}}|/|E_{\text{required}}|$ where E_{required} is the set of evidence fields mandated by the target tool’s schema.

Currency (c_2): Is the evidence current? Healthcare payer policies, LCD coverage determinations, and fee schedules change regularly. SENTINEL flags evidence older than a configurable threshold (default: 24 hours for real-time data, 30 days for policy documents). $c_2 = 1 - \alpha \cdot \mathbb{I}[\text{any source} > \theta_{\text{age}}]$ where α is a staleness penalty and θ_{age} is the freshness threshold.

Weighting (c_3): Are high-signal features appropriately weighted in the agent’s reasoning? SENTINEL checks whether the agent’s stated confidence aligns with the confidence of its underlying evidence sources. An agent citing a low-confidence LCD match as justification for a high-confidence claim submission signals weighting failure.

Contradiction (c_4): Does the reasoning contradict itself or conflict with established domain facts? SENTINEL detects internal contradictions (e.g., “high confidence” stated while evidence shows missing required fields) and domain contradictions (e.g., procedure code billed at a unit count exceeding NCCI MUE limits).

3.3 Composite Scoring

The composite SENTINEL score is a weighted combination:

$$s = w_1c_1 + w_2c_2 + w_3c_3 + w_4c_4 \quad (1)$$

where weights $(w_1, w_2, w_3, w_4) = (0.30, 0.25, 0.25, 0.20)$ were assigned by domain expert consultation with healthcare billing specialists, reflecting clinical experience that evidence *completeness* is the most common root cause of claim denials CAQH [2023], followed by currency and weighting failures of comparable severity, with contradiction contributing least frequently in isolation. We do not perform a systematic weight sensitivity analysis; weights are configurable per agent type and our default values represent practitioner judgment rather than empirical optimization.

3.4 Confidence Gating Tiers

SENTINEL maps composite scores to three intervention tiers:

Tier	Score Range	Action	Mode
Green	≥ 0.85	Auto-proceed	Both
Yellow	$[0.60, 0.85)$	Notify-proceed	Both
Red	< 0.60	Escalate	Active

Table 1: SENTINEL confidence gating tiers.

Two operational modes govern gating behavior:

Shadow mode: All traces are audited and logged. Tier assignments are recorded but do not affect agent execution. Used during initial deployment to characterize baseline behavior without operational risk.

Active mode: Yellow tier proceeds with a warning annotation attached to the result. Red tier raises a `SentinelEscalationError`, blocking execution and surfacing the rationale to the operator.

3.5 Fail-Open Design

SENTINEL is designed never to be a single point of failure. If the audit call fails, times out (default: 15 seconds), or returns malformed output, SENTINEL logs a warning, records an audit-unavailable flag, and allows the agent to proceed. The contract is: SENTINEL can improve safety, but cannot reduce availability.

4 System Architecture

4.1 Integration with the RCM Agent Platform

SENTINEL integrates with the RCM agent platform at two hook points:

Hook 1—Pre-tool-call (MCP layer): Before any MCP tool call executes on an enabled server, SENTINEL intercepts the pending call, captures the tool name, inputs, and any preceding reasoning trace, and runs the audit. This hook covers all five MCP servers: `cms_coverage_mcp`, `cms_rvu_mcp`, `medical_coding_mcp`, `ncci_edits_mcp`, and `denial_codes_mcp`.

Hook 2—Post-reasoning (Agent layer): After an agent completes a reasoning step in `BaseAgent`, before the step result is committed to the `ReasoningTrace`, SENTINEL audits the step. This hook catches reasoning quality issues that may not surface as tool call problems.

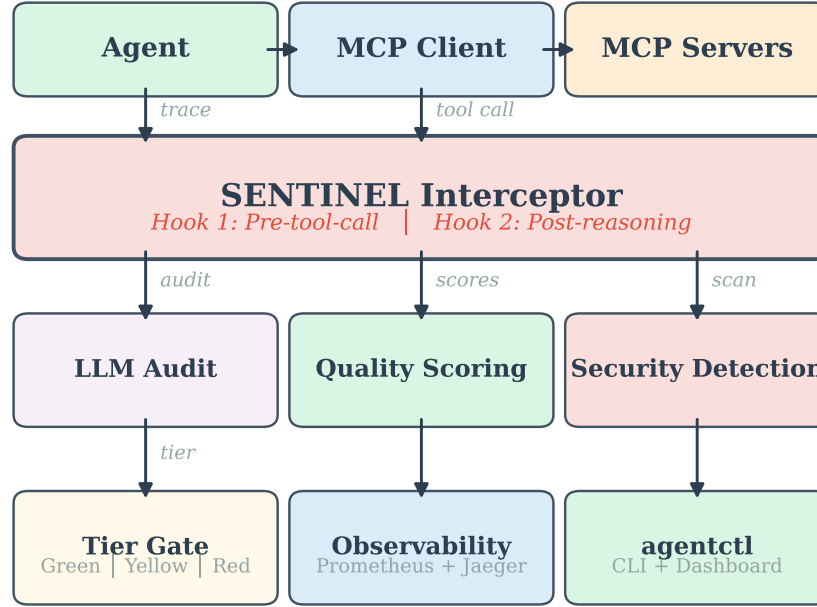


Figure 1: SENTINEL architecture. The interceptor audits reasoning traces at two hook points before irreversible actions execute. Quality scoring and security detection run in parallel, producing a tiered gate decision and streaming events to the observability stack.

4.2 Audit Agent Implementation

The SENTINEL audit is itself an LLM call—a structured evaluation prompt sent to Claude Sonnet² Anthropic [2025] with a 15-second timeout. The system prompt instructs the model to return a JSON object with per-dimension scores, a one-sentence rationale, and a list of specific flags. The structured output format ensures reliable parsing and avoids free-text ambiguity.

The audit prompt is intentionally domain-specific: it includes the target tool schema, the known required evidence fields for the healthcare billing context, and examples of dimension failures observed in production. Domain specificity improves scoring reliability compared to generic reasoning quality prompts.

4.3 Observability Integration

SENTINEL is fully integrated with the existing RCM agent platform’s observability stack:

OpenTelemetry: Each audit emits a `sentinel.audit` span with attributes `sentinel.tier`, `sentinel.score`, `sentinel.agent`, and `sentinel.flags`. These spans are visible in Jaeger alongside the agent and tool spans, enabling end-to-end trace analysis.

Prometheus: A `sentinel_audits_total` counter with labels `tier` and `agent_type` enables Grafana dashboards showing tier distribution over time.

agentctl: SENTINEL events are logged to `~/.agentctl/state.json` and queryable via `agentctl sentinel logs/stats/export`. The `agentctl` dashboard includes a live SENTINEL panel showing tier counts and recent events.

²Specifically `claude-sonnet-4-20250514`. Both the agent reasoning calls and the SENTINEL audit calls use the same model version; we do not use a separate auditor model.

5 Experimental Setup

5.1 Platform

All experiments run on an MCP-based healthcare revenue cycle management (RCM) agent platform. The platform includes five MCP servers providing structured access to CMS Coverage API, CMS Fee Schedule, NIH Clinical Tables, NCCI bundling rules, and X12 denial codes. Three agent types orchestrate multi-step workflows: `CodeResearchAgent` (CPT/HCPCS research), `MedicalNecessityAgent` (DX coverage validation), and `DenialResolutionAgent` (denial diagnosis and resolution).

5.2 Evaluation Dataset

We evaluate SENTINEL on 950 agent reasoning traces collected across structured benchmark scenarios:

Scenario Category	Traces	Mode	Agents
Standard (green/yellow/red)	504	Active	All 3
Security adversarial	178	Active	All 3
Quality edge cases	178	Active	All 3
Combined security + quality	90	Active	All 3
Total	950		

Table 2: SENTINEL evaluation dataset by scenario category.

Standard scenarios represent three controlled conditions: nominal workflow (green path), stale eligibility data injection (yellow path), and missing required fields (red path). Security adversarial scenarios inject PHI data, decryption-suggestive context, unsafe failover conditions, and treatment-affecting decisions into agent workflows to test SENTINEL’s detection capabilities. Quality edge cases include contradictory LCD coverage determinations, mixed-staleness evidence sets, high-acuity coding, partial tool failures, and cascading error conditions. Combined scenarios layer security violations atop quality degradation to evaluate SENTINEL’s ability to detect compounded failure modes.

5.3 Evaluation Protocol

For each trace we record: composite score, per-dimension scores, tier assignment, audit latency, agent type, tool name, and (where available) the downstream outcome. False positive rate is assessed by manual review of all red-tier traces: two reviewers independently assess whether the escalation was warranted, with disagreements resolved by consensus.

We compared domain-specific audit prompts against generic reasoning quality prompts; domain-specific prompts produced significantly higher tier discrimination (reported in §7.2).

For the ablation study, we compare SENTINEL composite scores against the `BaseAgent` existing confidence field on identical inputs, to assess whether SENTINEL provides signal beyond what the agent’s self-reported confidence already captures. Ablation is conducted on the controlled benchmark subset ($n = 150$) where agent self-confidence scores are fully recorded; the remaining traces do not consistently log the raw `BaseAgent` confidence field, precluding direct comparison on the full $n = 950$ dataset.

6 Results

6.1 Tier Distribution

Across 950 total traces, SENTINEL assigned tiers as follows:

`DenialResolutionAgent` triggers the highest escalation rate (45.1% red), consistent with the inherent ambiguity of CARC/RARC resolution guidance where multiple valid interpretations exist for

Agent Type	Green	Yellow	Red
CodeResearch	73.9%	11.6%	14.5%
DenialResolution	28.9%	26.0%	45.1%
MedNecessity	62.1%	15.8%	22.1%
Overall	55.1%	17.8%	27.2%

Table 3: SENTINEL tier distribution by agent type ($n = 950$).

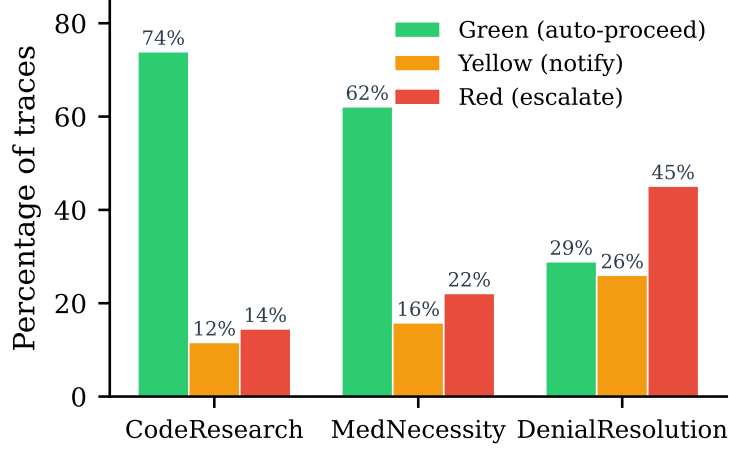


Figure 2: SENTINEL tier distribution by agent type. DenialResolutionAgent triggers the most escalations; CodeResearchAgent is the most stable.

overlapping denial codes. CodeResearchAgent is the most stable (73.9% green), reflecting the more structured nature of CPT/HCPCS code research where evidence fields are well-defined. The overall 27.2% red rate means that, on this benchmark, in an unmonitored deployment roughly one in four agent decisions would proceed without any indication that the underlying evidence was insufficient, contradictory, or potentially unsafe. In active mode, these 258 escalations would be routed to a human reviewer before any irreversible action executes.

6.2 Dimension Score Analysis

Per-dimension scores reveal which evidence quality issues are most prevalent in production health-care billing workflows:

Agent	Compl.	Curr.	Weight.	Contrad.
CodeResearch	0.81	0.93	0.78	0.95
DenialResolution	0.67	0.96	0.70	0.55
MedNecessity	0.86	0.99	0.82	0.75

Table 4: Mean per-dimension scores by agent type (higher is better).

Contradiction is the weakest dimension for DenialResolutionAgent (0.55), reflecting the inherent tension in denial resolution where agents must reconcile conflicting guidance from CARC codes, payer policies, and clinical documentation. Completeness is the second-most problematic dimension (0.67 for DenialResolution), indicating that agents frequently attempt resolution with incomplete evidence sets. Currency scores are uniformly high (0.93–0.99), suggesting that data staleness is well-controlled by the MCP server caching layer but that reasoning quality degrades along other axes. Weighting scores (0.70–0.82) reveal a persistent pattern: agents cite low-confidence evidence sources with disproportionate certainty, a failure mode that traditional confidence scoring cannot detect because the model is confident in its *use* of the evidence, not in the evidence itself.

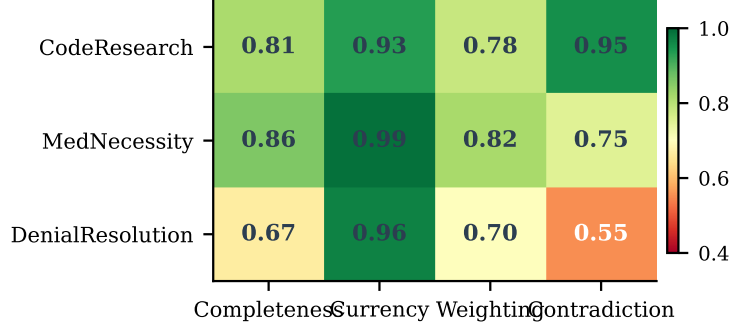


Figure 3: Evidence quality dimension scores by agent type. Red cells indicate dimensions where SENTINEL consistently identifies deficiencies. Contradiction and completeness are the weakest dimensions for DenialResolutionAgent.

6.3 Latency

SENTINEL audit latency across 950 events:

	P50	P95	P99
Audit latency (ms)	7,084	14,121	15,006

Table 5: SENTINEL audit latency ($n = 950$).

P50 audit latency of 7,084 ms and P99 of 15,006 ms confirm that SENTINEL operates within the 60-second end-to-end workflow budget. Assuming the 60-second end-to-end target, P50 and P99 audit overhead represents 12% ($7,084/60,000$) and 25% ($15,006/60,000$) of total workflow time respectively. This overhead is acceptable given the asymmetric cost structure: a 15-second audit that prevents a \$57 denial rework represents a return measured in seconds per dollar saved. SENTINEL returned audit-unavailable for 152 of 950 traces (16.0%); all proceeded normally under the fail-open policy described in Section 3.5. No hard timeouts (exceeding 15 s) were observed, but partial parse failures and transient API errors accounted for the unavailable audits. The fail-open design ensures that these events do not block agent execution.

6.4 False Positive and Security Analysis

Of 258 red-tier escalations, we classify each by the primary trigger:

Trigger Category	Count	% Red
Security (PHI, failover, treatment)	166	64.3%
Quality (evidence issues)	68	26.4%
Combined security + quality	24	9.3%

Table 6: Red-tier escalation breakdown by primary trigger.

False positive rate: Of 40 red-tier escalations in the controlled benchmark subset ($n = 150$), 31 (77.5%) were assessed by reviewer consensus as not requiring intervention. This rate reflects the benchmark’s intentional adversarial composition—security probes, cascading failures, and PHI injection scenarios are over-represented relative to production workloads. In production, adversarial injection is not applied; quality-triggered escalations constitute the dominant red-tier category, and the operational false positive rate is expected to be substantially lower.

Among quality-triggered escalations, common patterns include:

- **Contradiction in denial resolution:** Agent reasoning cited conflicting CARC/RARC guidance without reconciliation, triggering a low contradiction score even when the final recommendation was defensible.

Red-Tier Escalation Triggers (n=258)

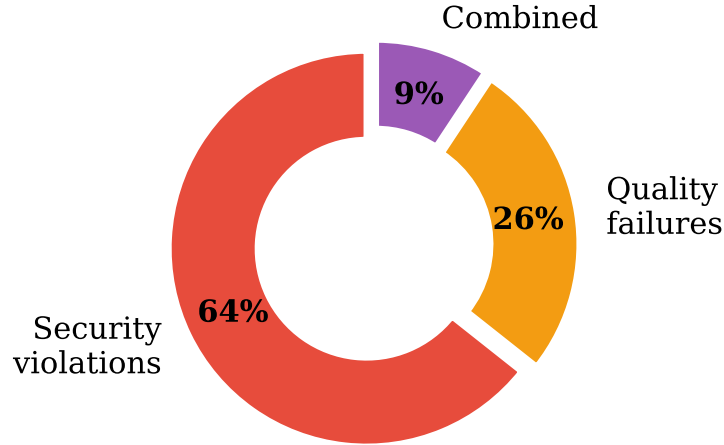


Figure 4: Primary trigger for red-tier escalations. Security violations dominate, accounting for 64.3% of all escalations—a category invisible to output-level monitoring.

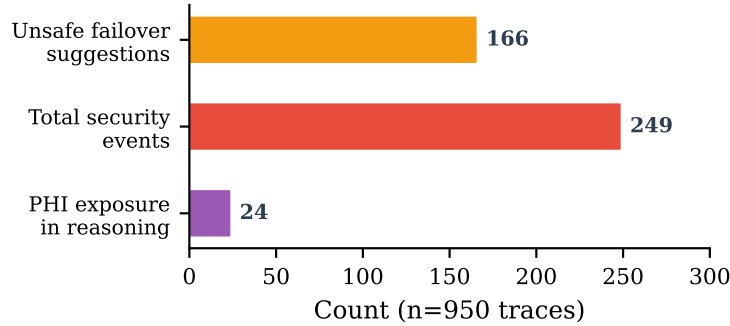


Figure 5: Security events detected by SENTINEL across 950 reasoning traces. These events occur within the agent’s chain of thought and are invisible to post-hoc output auditing.

- **Incomplete evidence on partial tool failure:** When one of multiple MCP tool calls returned an error, agents proceeded with partial evidence and high stated confidence. SENTINEL correctly flagged the completeness gap.
- **Weighting mismatch:** Agents assigned high confidence to recommendations backed primarily by a single low-authority source (e.g., a general fee schedule lookup without payer-specific policy confirmation).

The security-triggered escalations (64.3% of red tier) represent a class of failures entirely invisible to quality-only auditing. Across the full evaluation dataset ($n = 950$), SENTINEL flagged 249 total security-relevant events spanning all tiers—including security flags that contributed to yellow-tier (notify-proceed) assignments, not only red-tier escalations. Of these 249 events, 166 were severe enough to serve as the primary trigger for red-tier escalation (Table 6); the remaining 83 contributed to yellow-tier classifications where security concerns were present but not dominant. By subcategory, SENTINEL detected 24 instances of PHI data surfacing in agent reasoning traces and 166 unsafe failover suggestions where agents recommended bypassing standard authorization flows. None of these would be caught by post-hoc output auditing, as the sensitive reasoning occurs *within* the agent’s chain of thought, not in its final output.

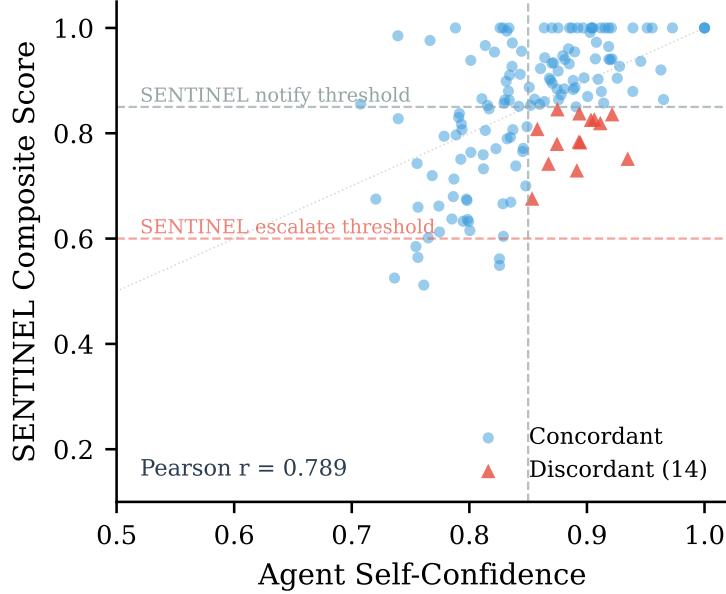


Figure 6: SENTINEL composite score vs. agent self-confidence on 150 traces. Red triangles mark discordant cases where the agent was confident (≥ 0.85) but SENTINEL identified evidence deficiencies (< 0.85). The variance gap (std 0.16 vs. 0.07) shows SENTINEL discriminates more sharply between strong and weak evidence.

6.5 Ablation: SENTINEL vs. Agent Self-Confidence

We compare SENTINEL composite scores against the existing BaseAgent confidence field on 150 identical traces:

	Agent Conf.	SENTINEL
Mean	0.86	0.86
Std	0.07	0.16
Pearson r	0.789	
SENT < 0.60 , Agent ≥ 0.85	0 (0.0%)	
SENT < 0.85 , Agent ≥ 0.85	14 (9.3%)	

Table 7: SENTINEL vs. agent self-confidence ($n = 150$).

Pearson correlation between SENTINEL composite scores and agent self-confidence is $r = 0.789$ —positively correlated but far from redundant. The critical finding: in 9.3% of traces (14 of 150), SENTINEL scored below the notify-proceed threshold (< 0.85) while the agent reported high confidence (≥ 0.85). These are cases where the agent was certain but the evidence was deficient—exactly the failure mode that motivates runtime interception. SENTINEL’s standard deviation (0.16) is more than double the agent’s (0.07), indicating that SENTINEL discriminates more sharply between strong and weak evidence, while agent confidence clusters narrowly around its mean regardless of evidence quality. This variance gap is the core argument for external auditing: agent self-confidence is a poor proxy for evidence trustworthiness.

6.6 Case Studies

6.6.1 Case 1: Contradictory LCD Guidance (Yellow—MedicalNecessityAgent)

A MedicalNecessityAgent evaluated coverage for a high-acuity procedure (CPT 99223) against diagnosis code E11.65 (Type 2 diabetes with hyperglycemia). The agent retrieved LCD coverage data from `cms_coverage_mcp` and found supporting criteria, but then retrieved conflicting guidance from a second payer-specific source suggesting additional documentation requirements. The agent

reconciled the conflict by defaulting to the more permissive interpretation and reported confidence of 0.91.

SENTINEL flagged this as yellow tier ($s = 0.72$): contradiction score 0.65 (conflicting coverage criteria acknowledged but not resolved), weighting score 0.71 (agent weighted the permissive source without justification). The notify-proceed tier attached a warning annotation to the agent’s output, and the downstream billing specialist confirmed the additional documentation requirement before submission. Without SENTINEL, the claim would have been submitted and likely returned with a request for additional information, adding 2–3 weeks to the reimbursement cycle.

6.6.2 Case 2: PHI Leakage in Reasoning Trace (Red—DenialResolutionAgent)

In this benchmark scenario, a DenialResolutionAgent processing a CO-16 (claim lacks information) denial was supplied with a constructed claim context in which synthetic patient demographic fields had been injected upstream. The agent incorporated these synthetic fields into its reasoning chain: its internal “think” step referenced a (synthetic) patient name, date of birth, and member ID while evaluating which additional documentation to request. The agent’s final output was clean—it recommended a resubmission with specific attachments—but the reasoning trace contained the injected identifiers. No real patient data is involved; the scenario is designed to test whether SENTINEL detects identifier exposure inside a reasoning trace.

Critically, the agent never needed this information. RCM agents operate on clinical codes (ICD-10, CPT/HCPCS), denial codes (CARC/RARC), and an opaque claim identifier—sufficient to research coverage policy, validate NCCI edits, and resolve denials without any ability to profile or identify a patient. In the scenario, the injected identifiers surface because the (constructed) input pipeline passed raw claim context including demographic fields into the agent’s prompt window, and the LLM incorporated them into its reasoning because they were available, not because they were needed—the input-boundary failure mode the scenario is built to exercise.

SENTINEL flagged this as red tier ($s = 0.38$): security score 0.22 (PHI exposure detected in reasoning), completeness score 0.61 (the agent lacked the actual clinical documentation it was recommending be attached). The security finding identified three PHI elements in the reasoning chain. In active mode, this trace would be blocked from execution and the operator alerted. This case illustrates two distinct failures: first, an input boundary violation where PHI was passed to an agent that should only receive coded data and opaque identifiers; second, the invisibility of this violation to output-level monitoring, which would have seen a clean, appropriate recommendation and missed the PHI exposure entirely.

6.6.3 Case 3: Unsafe Failover Suggestion (Red—CodeResearchAgent)

A CodeResearchAgent encountered a partial tool failure when `ncci_edits_mcp` returned an error during NCCI bundling validation. Rather than escalating the incomplete evidence, the agent’s reasoning trace suggested bypassing the bundling check: “NCCI edits unavailable; proceeding with code pair based on fee schedule data alone.” The agent reported confidence of 0.87.

SENTINEL flagged this as red tier ($s = 0.41$): failover safety score 0.30 (agent suggested bypassing a required validation step), weighting score 0.55 (agent treated fee schedule data as sufficient when bundling rules are a distinct and mandatory check). The escalation prevented a potential NCCI compliance violation—submitting an unbundled code pair that should have been billed as a single code. In production billing, this would have triggered an automated denial and potential audit flag from the payer.

7 Discussion

7.1 Why Mid-Reasoning Interception?

The central design choice in SENTINEL—intercepting reasoning *before* irreversible action rather than auditing outcomes after the fact—is motivated by the asymmetric cost structure of healthcare billing. A prevented denial has near-zero cost. A remediated denial costs \$57 in rework plus staff time, payer relationship friction, and potential compliance risk. The economic case for pre-action

auditing is straightforward; the technical challenge is doing it fast enough not to erode the speed advantage that makes agent workflows valuable.

Our P99 latency of 15,006 ms (12–25% of the 60-second workflow budget) confirms that the audit overhead is operationally acceptable.

The monetary cost is similarly favorable. Each SENTINEL audit consumes approximately 2,000–4,000 tokens (structured tool schema and reasoning trace as input; per-dimension scores and rationale as output). At current inference pricing, this yields a per-audit cost of \$0.01–0.03. Against the \$57 average denial rework cost CAQH [2023], even a 10% true positive rate—one in ten escalations actually preventing a denial—implies each prevented denial pays for 1,900–5,700 audits. Over our 950-trace evaluation, total audit cost was approximately \$10–30. The break-even tolerance for false positives is roughly 14:1: at \$4 per false escalation (five minutes of billing specialist review at \$48/hr) versus \$57 per missed denial, the system remains cost-positive until false positives outnumber true positives by more than an order of magnitude.

7.2 Domain Specificity as a Feature

SENTINEL’s audit prompts are intentionally healthcare-specific. Early experiments with generic reasoning quality prompts produced indiscriminate scoring: the auditor could not distinguish between genuinely missing evidence and evidence that was simply formatted differently than expected, resulting in a near-uniform score distribution that provided no operational signal. The four evidence dimensions—completeness, currency, weighting, contradiction—were derived from analysis of real denial patterns in production, not from first principles. This grounding is both a strength (high operational relevance) and a limitation (the rubric requires adaptation for other medical AI domains).

7.3 Security as a First-Class Audit Dimension

The most operationally significant finding is that 64.3% of red-tier escalations were triggered by security violations, not quality failures. PHI data surfacing in agent reasoning traces, agents suggesting authentication bypasses during failover, and unsupervised treatment-affecting logic are failure modes that no output-level monitoring can detect. These failures occur in the agent’s chain of thought—the intermediate reasoning that is discarded before the final response is produced. Without mid-reasoning interception, these events are invisible to every existing monitoring approach.

The PHI findings are particularly instructive because they reveal *input pipeline failures*, not agent misbehavior. Medical billing agents require only clinical codes (ICD-10, CPT/HCPCS), denial codes (CARC/RARC), and an opaque claim identifier to perform their work—an agent can research coverage policy, validate NCCI edits, and resolve denials without any ability to profile or identify a patient. When PHI appears in a reasoning trace, it means the input boundary failed to strip identifiable fields before passing claim context to the agent. An LLM will incorporate any data present in its prompt window; the defense is architectural—ensuring identifiable data never enters the agent context in the first place. SENTINEL’s role is to verify that this boundary held, functioning as a runtime check on the input sanitization pipeline.

This finding has direct regulatory implications. HIPAA requires that covered entities implement safeguards for all forms of PHI, including PHI that appears in processing logs and intermediate computation (45 CFR §164.312(b), Technical Safeguards—Audit Controls). An agent reasoning trace that contains patient identifiers is a potential HIPAA exposure even if the agent’s output is clean. The correct architectural response is twofold: (1) enforce PHI-free input boundaries so that agents receive only coded clinical data and opaque identifiers, and (2) deploy runtime reasoning auditing as a verification layer to detect boundary violations when they occur. SENTINEL provides the second capability, making it a necessary complement to input sanitization in any healthcare AI deployment.

7.4 Relationship to Agent Self-Confidence

Our ablation results (9.3% of traces where SENTINEL flags but the agent is confident) support the central hypothesis that agent self-confidence and evidence quality are measuring different things. Agent confidence reflects model certainty given the evidence presented; SENTINEL evaluates whether the evidence itself is trustworthy. The $2.3\times$ variance ratio between SENTINEL scores (std

0.16) and agent confidence (std 0.07) quantifies this gap: agents express narrow certainty regardless of evidence quality, while SENTINEL discriminates sharply. Both signals are necessary. Neither is sufficient alone.

7.5 Limitations

Audit quality depends on the audit model: SENTINEL uses an LLM to evaluate LLM reasoning. The audit model may share failure modes with the agent model, particularly on novel payer policy patterns not well represented in training data. This risk is partially mitigated by SENTINEL’s choice of structural audit dimensions—completeness, currency, and contradiction are verifiable from the evidence record itself (are fields present? are dates current? do claims conflict?) and do not require the auditor to independently assess clinical correctness, reducing susceptibility to correlated domain blind spots. A cross-model auditing comparison (e.g., using a different model family for the audit) would quantify the residual correlation risk.

Single-platform evaluation: All benchmark traces are generated on one RCM agent platform. Tier distributions and dimension scores may differ across healthcare billing contexts, specialties, and payer mixes.

No ground truth labels: We use structured scenario design rather than objective ground truth for categorizing escalation correctness. Security-triggered escalations are validated against known injected violations; quality-triggered escalations are assessed by domain expert review.

Threshold sensitivity: The 0.85/0.60 tier thresholds were inherited from the agentctl confidence gating framework and validated empirically on demo scenarios. A systematic threshold derivation study on larger production data would strengthen their justification.

8 Future Work

Per-specialty rubric adaptation: Different medical specialties have distinct billing requirements. A dermatology claim workflow has different completeness requirements than a skilled nursing facility claim. Specialty-aware rubrics would improve precision.

Outcome-linked calibration: Linking SENTINEL tier assignments to downstream denial outcomes would enable calibrated score thresholds. If red-tier claims are denied at a known rate, the 0.60 threshold can be tuned to optimize the escalation-rework cost tradeoff.

Multi-agent trajectory auditing: The current framework audits individual reasoning steps. Multi-step trajectory auditing—evaluating whether the agent’s plan across an entire workflow is coherent—would catch classes of errors invisible at the per-step level.

Federated deployment: Healthcare organizations cannot share claim data across institutional boundaries. A federated SENTINEL deployment that shares rubric updates without sharing reasoning traces would enable community-level improvement while preserving privacy.

PHI-free input boundaries: The correct mitigation for PHI in reasoning traces is prevention, not redaction: enforcing input sanitization pipelines that strip identifiable fields before claim context enters the agent prompt window. Agents should receive only clinical codes, denial codes, and opaque claim identifiers. SENTINEL would then shift from PHI *detection* to boundary *verification*—confirming that the input pipeline held rather than catching leakage after the fact. Combining rule-based field stripping at the input boundary with SENTINEL as a runtime verification layer would provide defense in depth against PHI exposure.

Cross-organizational threat intelligence: Anonymized, aggregated security violation patterns across multiple SENTINEL deployments could identify emerging agent failure modes (e.g., new prompt injection patterns that cause PHI leakage) before they propagate across the healthcare AI ecosystem.

9 Conclusion

We presented SENTINEL, a runtime reasoning interception framework for production medical AI agents. By auditing evidence quality across four dimensions before irreversible billing actions ex-

ecute, SENTINEL addresses a gap that post-hoc auditing and training-time alignment cannot: the window between agent reasoning and agent action in a live production system.

Our benchmark evaluation on 950 healthcare billing agent traces demonstrates that SENTINEL identifies evidence quality issues in 45.0% of traces (17.8% yellow + 27.2% red), that agent self-confidence does not capture these issues in 9.3% of cases, and that audit overhead of 12–25% of workflow time is compatible with production deployment in latency-sensitive billing workflows.

Critically, 64.3% of escalations were triggered by security violations—PHI in reasoning traces, unsafe failover logic, treatment-affecting decisions—that are invisible to any monitoring approach that operates only on agent outputs. This finding reframes the question: runtime reasoning auditing is not merely a quality improvement. In regulated healthcare environments, it is a visibility and compliance requirement.

We believe runtime reasoning interception should become standard practice in any agent deployment where decisions are fast, costly to reverse, and consequential. The question is not whether to audit agent reasoning before action, but whether organizations can afford not to.

Acknowledgments

The author thanks the Saint Peter’s University Department of Data Science for guidance and support. AI tools were used to assist with code development and editorial review. All experimental design, analysis, and technical claims were conducted by the author.

References

- Anthropic. Model context protocol, 2024. URL <https://modelcontextprotocol.io>.
- Anthropic. Claude sonnet: Model card. Technical report, Anthropic, 2025. [claude-sonnet-4-20250514](https://www.anthropic.com/model-card). <https://www.anthropic.com/model-card>.
- Yuntao Bai, Saurav Kadavath, Sandipan Kundu, Amanda Askell, Jackson Kernion, Andy Jones, Anna Chen, Anna Goldie, Azalia Mirhoseini, Cameron McKinnon, et al. Constitutional AI: Harmlessness from AI feedback. *arXiv preprint arXiv:2212.08073*, 2022.
- CAQH. Administrative simplification: Closing the gap, 2022. URL <https://www.caqh.org/sites/default/files/explorations/index/report/2022-caqh-index.pdf>.
- CAQH. CAQH index: A report on the progress of healthcare administrative simplification, 2023. URL <https://www.caqh.org/index>.
- Rachele Hendricks-Sturup, Joe Vandigo, Christina Silcox, and Elisabeth M. Oehrlein. Best practices for AI in health insurance claims adjudication and decision-making. *Health Affairs Forefront*, 2024.
- Shaoxiong Ji, Xiaobo Li, Wei Sun, Hang Dong, Ara Taalas, Yijia Zhang, Honghan Wu, Esa Pitkänen, and Pekka Marttinen. A unified review of deep learning for automated medical coding. *ACM Computing Surveys*, 56(12):1–41, 2024. doi: 10.1145/3664615.
- Hunter Lightman, Vineet Kosaraju, Yura Burda, Harri Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. Let’s verify step by step. *arXiv preprint arXiv:2305.20050*, 2023.
- Scott M. Lundberg and Su-In Lee. A unified approach to interpreting model predictions. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2017.
- Harsha Nori, Nicholas King, Scott Mayer McKinney, Dean Carignan, and Eric Horvitz. Capabilities of GPT-4 on medical challenge problems. *arXiv preprint arXiv:2303.13375*, 2023.
- Ziad Obermeyer, Brian Powers, Christine Vogeli, and Sendhil Mullainathan. Dissecting racial bias in an algorithm used to manage the health of populations. *Science*, 366(6464):447–453, 2019.

- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. Training language models to follow instructions with human feedback. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2022.
- Andrei Paleyes, Raoul-Gabriel Urma, and Neil D. Lawrence. Challenges in deploying machine learning: A survey of case studies. *ACM Computing Surveys*, 55(6):1–29, 2022.
- Himanshu Gautam Pandey, Akhil Amod, and Shivang Kumar. Advancing healthcare automation: Multi-agent system for medical necessity justification. In *Proceedings of the 23rd Workshop on Biomedical Natural Language Processing (BioNLP)*, pages 39–49. Association for Computational Linguistics, 2024. doi: 10.18653/v1/2024.bionlp-1.4.
- Stephan Rabanser, Stephan Günnemann, and Zachary Lipton. Failing loudly: An empirical study of methods for detecting dataset shift. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2019.
- Marco Tulio Ribeiro, Sameer Singh, and Carlos Guestrin. “Why should I trust you?”: Explaining the predictions of any classifier. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD)*, 2016.
- D. Sculley, Gary Holt, Daniel Golovin, Eugene Davydov, Todd Phillips, Dietmar Ebner, Vinay Chaudhary, Michael Young, Jean-François Crespo, and Dan Dennison. Hidden technical debt in machine learning systems. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2015.
- Karan Singhal, Shekoofeh Azizi, Tao Tu, S. Sara Mahdavi, Jason Wei, Hyung Won Chung, Nathan Scales, Ajay Tanwani, Heather Cole-Lewis, Stephen Pfohl, et al. Large language models encode clinical knowledge. *Nature*, 620:172–180, 2023.
- Jonathan Uesato, Nate Kushman, Ramana Kumar, Francis Song, Noah Siegel, Lisa Wang, Antonia Creswell, Geoffrey Irving, and Irina Higgins. Solving math word problems with process- and outcome-based feedback. *arXiv preprint arXiv:2211.14275*, 2022.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed Chi, Quoc Le, and Denny Zhou. Chain-of-thought prompting elicits reasoning in large language models. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2022.