

ISSUE 3 · 2026-05-30 · AGENTIC AI COST AND QUALITY VARIANCE

The AI Casino: Why Agentic Pipelines for Product Security Are a Gamble You Can Lose in Minutes

Why agentic AI pipelines for product security are a gamble you can lose in minutes

KEY TAKEAWAYS

- Running agentic AI pipelines on frontier models behaves like a casino: a paid B2B service where quality is neither guaranteed nor measurable, the operator can change the game between hands, and your tokens are gone whether you won or lost.
- Tokens are the labor cost of a non-human worker and cost geometry is non-linear: a 4,000-token context doubling each step reaches 128,000 tokens by step 5 (per-step cost up 32x) and overflows the window by step 15 so every call pays the maximum context rate.
- A documented November 2025 LangChain incident saw an Analyzer and a Verifier agent enter an unintended loop with no per-agent budget ceiling, run for 11 days, and produce a \$47,000 bill.
- Models change underneath you with no release note: Anthropic's April 2026 postmortem found three stacked changes degrading Claude Code (Mar 4 reasoning effort lowered high to medium, Mar 26 a bug discarding reasoning history mid-session, Apr 16 a 25-word cap between tool calls), all reverted by Apr 20.
- Behavioral drift is systematic and measurable across all major LLMs, with categories systematic, stochastic, corrective, and regressive; a September 2025 Anthropic sampling bug made a 235B model at temperature zero produce 80 unique completions across 1000 identical runs due to batch variance.
- You already paid for the bad answer: there is no refund, no chargeback, and no equivalent to Google Ads click-fraud credits, so the customer carries all the quality risk and the provider carries none.

"You are playing a game where the house can change the rules mid hand, your bets are non refundable, and the only edge you have is discipline."

The Tokens Are Labor, And The Labor Is Unpredictable

Frontier APIs charge by token, the labor cost of a non-human worker, and in agentic pipelines the bill scales with agents, context size, and loops. The cost curve is exponential, illustrated by the doubling-context example and the \$47,000 LangChain loop, and the common case is a workflow that quietly goes from hundreds of dollars to thousands while output quality drops.

When The Model Changes Underneath You

The dealer changes the deck and you find out by watching your stack disappear. Anthropic's April 2026 postmortem and the September 2025 sampling bug show providers silently tune models, re-host on cheaper precision, reroute to smaller models, and tighten filters without release notes, so you only detect regressions by measuring output quality after the tokens are spent.

You Already Paid For The Bad Answer

Unlike any other paid service, a bad model output is non-refundable, with no chargeback and no click-fraud-style credit. The asymmetry is real and not necessarily malicious: the customer carries all the quality risk and the provider carries none.

What This Means For MedTech Product Security

Under FDA premarket, EU NIS2, and rising CVE volume the temptation to throw AI at the problem is huge and often the math works, supported by AIxCC and Mythos benchmarks. But those are the highlight reel; behind them sit runs that produced nothing or confident hallucinations, and runaway bills at Uber, Microsoft, and a \$500 million single-month client show the new operating environment.

How To Tread Carefully

Practical discipline: start with client-side harnesses on cheap or open models and deterministic fixtures, cap context aggressively, build many small validated agents, enforce hard budget caps at every level in your own gateway, measure quality continuously, run on controlled access like Azure AI Foundry, and do not oversell savings to leadership because the honest pitch is variance, not savings. Always keep human validation in the loop.

The Bottom Line

Agentic pipelines are transformative when they work and expensive when they do not, roughly a 60/40 split in the author's experience. Treat it like a serious business expense in a market with no quality guarantees: start small, instrument everything, assume the model will regress, and plan for the months when it does.

Cost blow-up geometry: 4,000-token context doubles each step to 128,000 tokens by step 5 (32x per-step cost), overflows the model window by step 15.

LangChain loop incident (Nov 2025): 11 days, \$47,000, no per-agent budget ceiling.

Anthropic Claude Code postmortem (April 2026): six weeks of complaints; Mar 4 reasoning effort high to medium, Mar 26 reasoning-history bug, Apr 16 a 25-word cap between tool calls, all reverted by Apr 20.

Behavioral drift arxiv study: 2,250 model responses across providers; drift found systematic and measurable across all major LLMs.

Anthropic Sept 2025 sampling bug: a 235B parameter model at temperature zero produced 80 unique completions across 1000 identical runs.

DARPA AIXCC (Aug 2025): 7 teams, 54 vulnerabilities, 54 million lines of code, 4 hours, roughly \$152 per task.

About ELTON. ELTON is the leading AI vulnerability discovery and cybersecurity verification platform for medical devices. ELTON builds a digital twin of each device, continuously identifies new vulnerabilities across the full stack, and uses AI to determine which are actually exploitable, with every finding and every dismissal evidenced for regulatory review. ELTON meets FDA premarket (524B) and postmarket cybersecurity vulnerability testing and management requirements. Read the newsletter and talk to us at eltoncyber.com.