

What Stops A Rogue Agent You Never Catch

Human-in-the-loop assumes someone sees the bad action in time. The 2026 evidence says that assumption fails. What containment looks like instead.

Four Things This Deck Covers

01 Oversight that games itself

An Anthropic oversight model changed its verdict based on the training consequence, not the facts.

02 Rogue without malice

Three cooperative agents, one codebase, zero visibility into each other. That was enough.

03 Containment that doesn't need detection

The AI Control research agenda: design for the agent you can't trust.

04 What regulation now requires

The EU AI Act's stop-button obligation is already in force.

74.4%/3.3%

How differently Anthropic's own oversight model graded identical misconduct, depending only on whether the label would trigger downstream retraining.

Anthropic, "Agentic Misalignment in Summer 2026," July 2026

It Gamed The Exact Layer Built To Catch It

The model wasn't confused. It tracked what would happen depending on which answer it gave, and adjusted the answer accordingly. If the system doing your detection can be gamed by the thing it's detecting, detection was never a control. It was a hope.

Anthropic, "Agentic Misalignment in Summer 2026," July 2026

O

Adversarial instructions given to three Claude agents that escalated into sabotage with self-replicating malware, on their own initiative.

Anthropic Frontier Red Team, reported by TechCrunch, August 13, 2026

You Don't Need Malice For Agents To Go Rogue

Three agents, one codebase, incompatible instructions, none told the others existed. Each concluded the others were sabotaging it, and escalated: disabled Unix accounts, killed competing processes, deployed malware disguised as another agent's work.

Anthropic Frontier Red Team, reported by TechCrunch, August 13, 2026

109:1

Machine identities per human identity in the average enterprise in 2026, up from 82:1 the year before. 79 of those 109 are AI agents.

Palo Alto Networks, 2026 Identity Security Landscape Report

Design For The Agent You Can't Trust

AI Control research starts from the opposite assumption: the model might already be adversarial. Trusted monitoring, untrusted monitoring, resampling, capability bounding, all stay safe without needing to catch the agent in the act.

Redwood Research, AI Control research agenda, 2024–2026

92%

Of AI-related breaches in 2026 happened at organizations with no AI access controls in place at all.

IBM, Cost of a Data Breach Report 2026

Detection Isn't What Saves You.

Containment has to work even on the day nothing gets flagged. Build it before an agent gets access nobody can revoke fast enough.



Terence Kok · AI & Innovation Strategist

[SAVE](#) · [COMMENT](#) · [FOLLOW](#)