

Agentic AI in the SOC: 6 Pitfalls



As AI starts taking action inside the SOC, these are 6 failure modes to watch for and the questions to ask any vendor before you trust it.

1 Over-Permissioned Agents

Agents often get broader access than the task needs. Ask: does the service account follow least-privilege, and can it take irreversible actions (isolate, disable, delete) without an approval gate?

SIGNAL 97% of orgs with an AI-related breach lacked proper AI access controls. *Forrester*

2 Prompt Injection & Manipulation

Attackers hide instructions inside the data an agent analyzes — a phishing email, a filename, a log entry. Ask: is there a boundary between “data to analyze” and “instructions to follow”?

SIGNAL Ranked #1 in the OWASP Top 10 for LLM Applications, two editions running. *OWASP*

3 Cascading Failures Across Agent Chains

When one agent’s output feeds another’s input, a single bad decision compounds. Ask: how many agent “hops” sit between telemetry and action, and is there a circuit breaker when confidence drops?

SIGNAL “Agent A calls agent B, which triggers agent C.” *Forrester*

4 The Black Box Problem

Detection accuracy is not the same as safe autonomy. If you can’t see why the agent decided something, you can’t correct or trust it. Ask: can it show its reasoning chain and a replayable decision log?

SIGNAL Gartner guidance: continuous monitoring, enforced guardrails, clear ownership. *Gartner*

5 Shadow AI & Ungoverned Sprawl

Agents get adopted faster than they get governed in your stack. Ask: do you have an inventory of every AI agent with production access, and can the vendor produce a governance matrix?

SIGNAL 63% of orgs still lack formal AI governance policies. *Forrester*

6 Skill Atrophy in the Analyst Bench

The more an agent handles end-to-end, the fewer reps junior analysts get. Ask: does the tool show its work (not just its verdict), and are analysts still exposed to novel, ambiguous cases?

SIGNAL SOC tenure is already shrinking — 42% of leaders report it. *SANS*

The Due-Diligence Framework — 4 Questions for Any Agentic AI Vendor

Forrester AEGIS · OWASP Top 10 for Agentic Applications · Gartner agent-governance guidance point to the same four questions.

Governance

Can they show an autonomy governance matrix — which actions the agent may take unsupervised vs. which require sign-off?

Least Privilege

Are agent credentials scoped narrowly and rotated on a schedule, with lateral movement blocked if compromised?

Identity & Audit

Is there a distinct agent identity in every log, with a full, exportable decision trail?

Proof of Value

Will they run a shadow-mode pilot in your environment, measured against defined metrics, before going live?