

The Buyer's Guide for  
the AI SOC Skeptic

# TRUST ISSUES



# Index

|                                       |    |
|---------------------------------------|----|
| The Rise of the AI SOC                | 03 |
| <hr/>                                 |    |
| The Trust Gap                         | 03 |
| <hr/>                                 |    |
| Where AI Goes Wrong in the SOC        | 04 |
| <hr/>                                 |    |
| What Makes an AI SOC Trustworthy?     | 05 |
| <hr/>                                 |    |
| Ten Questions for Every AI SOC Vendor | 06 |
| <hr/>                                 |    |

# The Rise of the AI SOC

Traditional SOAR was built for a slower world. Playbooks and manually-authored rules can't keep pace with threats that evolve in real time or with the sheer volume of change AI has unleashed across the enterprise.

In response, a new category has emerged: the AI SOC. Rather than relying on static automation, AI SOC's promise something fundamentally different: security operations that reason, adapt, and act with the same speed and flexibility as the risks they're meant to counter.

It's a compelling vision. But it also raises a hard question: can you trust AI in your SOC?

## The Trust Gap

That question sits at the center of a growing tension in the industry. Security teams need the speed of AI. Attacks now move faster than any human analyst can, and the volume of alerts, code, and infrastructure being generated daily has outpaced manual review entirely. But speed alone isn't the goal; trust is the prerequisite for using that speed at all.

This is the real trust gap: not a question of whether AI can act fast enough to defend and respond, but whether security teams can trust it to act correctly when it does. Every autonomous decision an AI SOC makes — to escalate, to contain, to dismiss — carries consequences, and for a SOC analyst who may need to defend that decision to a CISO, an auditor, or a regulator, "the AI said so" isn't an answer. It's a liability. The tension, then, isn't AI versus human. It's velocity versus verification. Closing that gap is what determines whether AI SOC becomes a force multiplier or just a faster way to be wrong.

---

*79% of SOC's are already using AI or ML tools, but only 36% have built that into a defined, governed workflow—2026 SANS SOC Survey Insights: A Decade of Evolution in Cyber Defense, June 2026*

**How to use this guide:** In this guide, we'll explain where trust breaks down across the AI SOC market and give you specific questions to ask vendors as you evaluate solutions for your security stack.

# Where AI Goes Wrong in the SOC

Most AI SOC products put one frontier model in charge of everything — triage, verdicts, response. Four places that chain snaps.

## Monolith agents, delegated authority

One top-level agent spawns sub-agents and passes credentials down the chain.

### TRUST ISSUE:

When something acts, nothing proves which agent authorized what.

## Guardrails that are just prompts

"The agent is instructed not to" is a request, not a control.

### TRUST ISSUE:

A prompt guardrail can be reasoned around by the very model it's meant to constrain.

## Frontier reasoning on every step

Agentic workloads consume  $\sim 1,000\times$  the tokens of a chatbot session.

### TRUST ISSUE:

Run a frontier model on lookup-table problems and the economics collapse at exactly the alert volume where you need it most.

## The agent grades its own homework

One model investigates, decides, and confirms its own verdict.

### TRUST ISSUE:

Unchallenged verdicts ship hallucinations to your response playbooks.

# What Makes an AI SOC Trustworthy?

01

## Agent Harness

Scoped tools, brokered credentials, bounded knowledge — hard bounds sit outside the model's reasoning.

02

## Challenger Pattern

A proposer builds the verdict, a challenger attacks it, an adjudicator decides.

03

## Zero Credential Exposure

Credentials belong to audited workflows. The agent invokes an ability, never a secret.

04

## Ability, Not Authority

Every action is a granted ability. Irreversible actions gate on approval by default.

## THE TRUST TEST

# Ten Questions for Every AI SOC Vendor

01

Which steps run as probabilistic AI reasoning, and which as deterministic workflows?

✓ a mix of both deterministic workflows for repeatable stages and reasoning only where judgment pays.

✗ "the agent handles everything."

02

Are agent bounds enforced by the platform, or requested in a prompt?

✓ a hard-coded harness.

✗ "the agent is instructed not to..."

03

What physically stops a runaway agent mid-execution?

✓ circuit breakers, resource limits, isolated execution.

✗ no concrete mechanism.

04

Does any second system challenge a verdict before action is taken?

✓ a separate challenger and adjudicator.

✗ one model confirms itself.

05

What happens when the agent is uncertain? Is "inconclusive" a possible verdict?

✓ inconclusive closes the case and feeds detection engineering.

✗ it always forces a verdict.

06

**Does the agent hold credentials? Can it spawn sub-agents that inherit them?**

✓ the agent never holds a secret; owned workflows do; no sub-agents.

✗ agent has direct access to credentials.

07

**Can autonomy vary per action and per use case, with approval gates on irreversible steps?**

✓ autonomy is a per-action dial; irreversible actions gate on approval.

✗ all-or-nothing autonomy.

08

**Can my team inspect every decision, input, and action in one auditable record?**

✓ one immutable audit trail.

✗ partial logs scattered across tools.

09

**Does it execute remediation end to end, or only recommend it?**

✓ closes cases with scoped, reversible workflows.

✗ it generates tickets for humans to clean up.

10

**How does cost scale with alert volume when reasoning runs on every step?**

✓ cheapest tier per stage; reasoning only where it pays.

✗ frontier reasoning on every step (~1,000× tokens).



# See BlinkOps in Action

The Agentic Security Operations Platform  
that goes beyond AI SOC.

[blinkops.com](https://blinkops.com)