

WHITE PAPER

AI + Network Control

Visibility, Risk Prioritization,
and Automation in the
Age of Agentic NetOps

Written by **Matt Bromiley**
July 2026



Introduction

McKinsey's latest *State of AI* survey found that 88% of organizations now report regular AI use in at least one business function, compared with 78% the previous year.¹ We're past the point of asking "Will you adopt AI?" Instead, we're examining how that adoption is being handled. Is all that AI being implemented correctly? Is it inventoried, governed, observed, and controlled with the same rigor expected of any other system that touches production?

For network and security operations, that question is about to get harder. The next step of enterprise AI is agentic systems that pursue goals and select their own actions rather than executing prescribed steps. When applied to the network, agentic NetOps means software that can correlate a change, model its impact, and propose or even execute a remediation without a human choosing each step. Security teams read that sentence and see risk before they see efficiency. Remove the human from the loop, and three exposures show up immediately:

- A lack of guardrails on *what* the system can do
- Data leakage as the agent reaches across systems it was *never* scoped to touch
- Accountability that gets harder to assign when no person chose the action

Not every organization is running agents today. Most that are keep them narrow and single-tracked. The direction of investment is unambiguous; however, the operating model in most enterprises has not kept pace with today's modern AI footprint, let alone an autonomous one.

This paper is written for the security and NetOps practitioners who will own that implementation, because they are the ones who get paged when it goes wrong. We define what NetOps means and the areas of implementation—governance, visibility, and control—that close the operational maturity gaps those implementations are opening.

We encourage you to read this paper against your own environment. Where does your agentic or autonomous NetOps implementation actually stand, how far along the timeline are you, and have you already run into the problems this paper describes? If you have, you are not alone, and you are not late—but the window to get ahead is now.

¹ McKinsey & Company, "The state of AI in 2025: Agents, innovation, and transformation," November 2025, www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai

What Agentic NetOps Actually Means

Our working definition for agentic NetOps is as follows: a system pursuing goals or intents and selecting actions, where operators previously chose the actions and the system executed them. Operators set policy and outcomes, and the system reasons against the current state to decide what to do.

Bain's 2025 Technology Report describes four levels of agentic AI deployment:

1. Information-retrieval agents that summarize and recall
2. Single-task agents that complete one workflow
3. Cross-system orchestration agents that span multiple tools
4. Multi-agent constellations that coordinate agents²

Bain places the current center of investment and deployment velocity at levels two and three—single-task agents and cross-system orchestration (see Figure 1). Thus, most real agentic NetOps work happening today lives at these levels. Multi-agent constellations remain largely ahead of the enterprise and, in our experience, are also new or untested within security technologies.

It is worth noting that the industry has not settled on a single definition yet. The term “agentic” is applied to everything from a scripted chatbot to a fully autonomous remediation engine. At this stage, it is difficult to cut cleanly between what marketing claims and what the technology does in production. We use this definition throughout this paper so the rest of the argument has firm ground, while acknowledging that vendors may define the term differently.

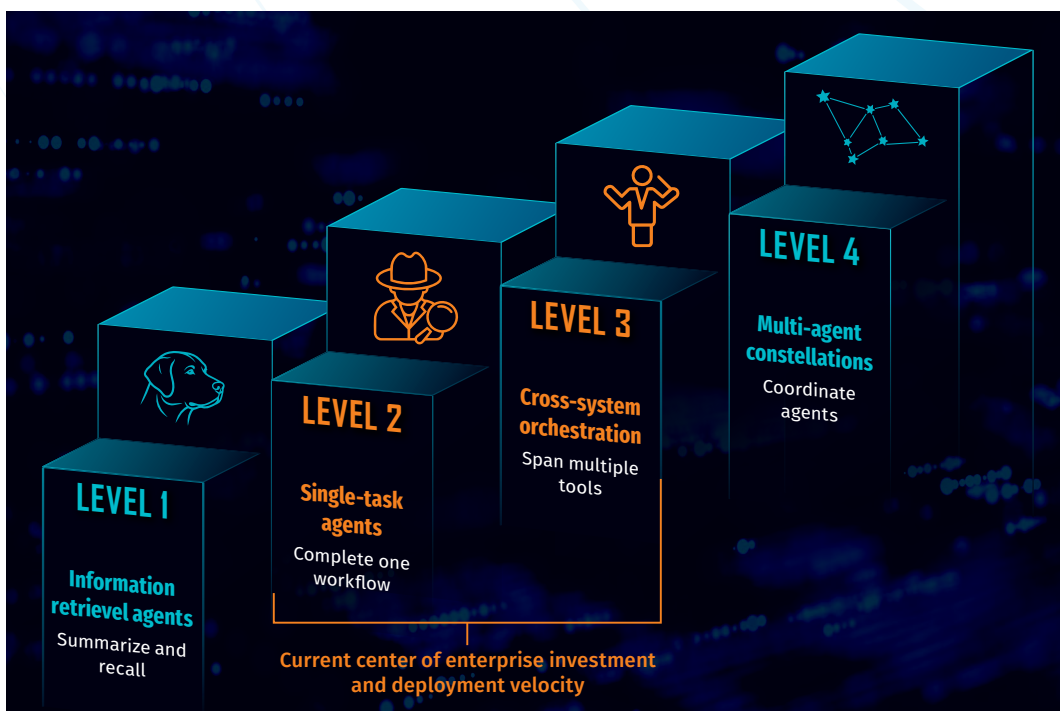


Figure 1. The four levels of agentic AI deployment described in Bain's 2025 Technology Report. Current enterprise investment and deployment velocity center on levels 2 and 3.

² Bain & Company, "Technology Report 2025," www.bain.com/insights/topics/technology-report/

As we researched industry resources for this paper, we identified a need for clarification: Three adjacent categories get conflated with agentic NetOps. The distinctions matter because each has its own tooling, telemetry, and policy model(s) that do not map cleanly onto autonomous agents:

- **AIOps** is anomaly detection and event correlation. It surfaces problems for humans to act on. Cisco's framing of agentic operations is that AIOps helps you see problems,³ while agentic NetOps helps you solve them and act.
- **Intent-based networking** translates declarative configuration into device-level state. It is closer to compilation than to reasoning. Agentic NetOps may use intent-based networking as a downstream effector, but the agent's job is to choose the intent based on current conditions, not to render the configuration.
- **Security orchestration, automation, and response (SOAR)** platforms run deterministic playbooks. The branching logic was written in advance by humans. Agentic NetOps assembles its own sequence of actions to pursue an outcome. The autonomy is in the planning *and* execution.

Where Adoption Actually Sits Today

According to McKinsey, 23% of organizations are scaling at least one agentic system somewhere, which may sound like a healthy adoption percentage. However, a granular look at the statistics reveals that no business function in the data reports more than roughly 10% of organizations with fully scaled agents in that function.

Or, better put: *AI deployments exist, but are narrow.*

Single-task or single-domain—consistent with Bain's level two and three observations—likely account for the majority of these agents. For network and security practitioners, the practical approach is to get governance, visibility, and control now, before they get outpaced by deployment (as we saw with cloud rollouts in many organizations).

**The data shows what we know:
Deployments are moving
quickly, but operational maturity
implantation needs to pick up the
pace. There is a window—albeit
one that is shrinking—ideal for
implementing the correct controls.**

³ Cisco, "What is agentic operations (AgenticOps)?" www.cisco.com/site/us/en/learn/topics/artificial-intelligence/what-is-agentic-operations-agenticops.html

Case Study: Operations as Normal?

Consider a common request that currently crosses several teams and tools: A user wants to open a connectivity path so an application can reach new or additional data sources. When handled manually, that request touches firewall management, a vulnerability scanner, a risk-scoring system, the device itself, and, finally, a compliance or audit record. In many organizations, this is often across multiple vendors and owners.

An agentic NetOps system collapses that sequence into *one* governed workflow (see Figure 2):

- It *receives the request*, capturing the requested connectivity between an application and a new data source along with the business context needed to evaluate and process the change.
- It *maps the path*, identifying what already exists along the request route—the existing rules, the devices, and the segments the new flow would traverse.
- It *assesses the risk*, pulling the organization’s risk profile for those segments and enumerating the vulnerabilities the new path would expose.
- It *implements the change* on the device itself, scoped to what was actually requested rather than opening the path wider than asked.
- It *validates that the change behaves* as intended and didn’t introduce any unintended flows or policy conflicts.
- It *documents the result(s)*, generating the compliance record: what changed, why, what was checked, and the evidence trail for audit.

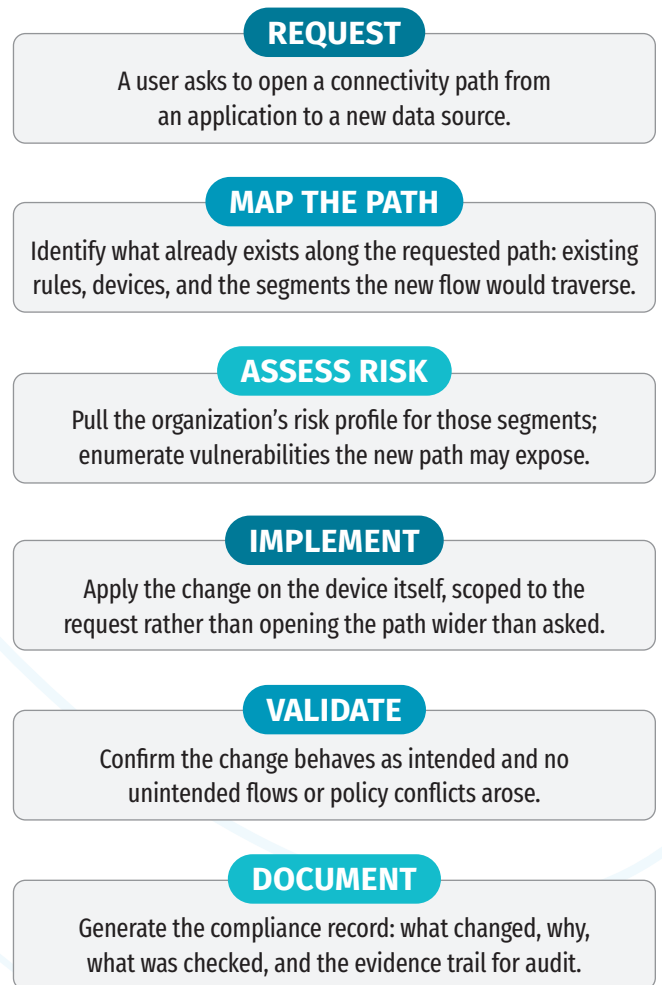


Figure 2. An Agentic NetOps Workflow for a Connectivity-Path Request

The capability that matters here is the agent mapping, automating, and then producing compliance documentation in place of a multi-system, multi-vendor sequence of manual controls. That is the class of problem agentic NetOps is being built to solve, and it is close enough to current enterprise practice to be a useful test of whether a given deployment is ready.

Picture this in a regulated environment with a quarterly audit cycle. The same request that used to generate a ticket, a scan, a risk sign-off, and a manual evidence package now produces all of it in one pass, with the documentation generated as a byproduct rather than reconstructed weeks later. That is the upside. The risk is what happens when the agent does all of that based on a policy that no longer reflects the organization’s intent.

Beyond Adoption: The Operational Maturity Gap

The gap between AI in production and the controls that govern it are the real story of AI in 2026—operational enterprise maturity.

According to McKinsey (the survey we're using as a baseline), 88% of organizations regularly use AI, up from 78% a year ago. Roughly one-third are scaling AI enterprise-wide, but only 7% describe themselves as fully scaled. Only 39% report any measurable earnings impact before interest and taxes (EBIT) from AI—and most of that attribution is under 5%.

Simply put: *The technology has arrived ahead of the operating model that can predictably extract value from it.*

Bain's *2025 Technology Report* helps explain the financial stakes. They project that 5–10% of total technology spending will flow toward foundational agentic capabilities over the next three to five years, with the share rising each year.

The infrastructure layer beneath this is being built now. Bain documents that Model Context Protocol (MCP) server counts grew sevenfold over the course of 2025.⁴ Agent-to-tool plumbing is going up faster than most operations teams are aware of or can even keep track of. The economic urgency is real, with an appreciable gap between the projection and what most enterprise environments can instrument and control today.

A practitioner should weigh any AI capability using the same method they weigh any other control: First, question its validity, then, measure the value it returns against the risk it introduces.

The typical enterprise has AI running in production, despite the absence of four things:

1. A systematic inventory of where AI is deployed
2. AI-aware policy enforcement on the network paths it uses
3. In-flow visibility into agent behavior
4. Audit trails sufficient to satisfy regulators and/or upcoming AI compliance

In Cisco's *2025 Cybersecurity Readiness Index*, which polled 8,000 security and business leaders across 30 markets, only 4% of organizations achieved its "Mature" readiness tier—the same number as the year before.⁵ Seventy percent remain in the bottom two stages; and 86% of leaders reported at least one AI-related security incident in the prior 12 months.

IBM's *2025 Cost of a Data Breach Report* identifies that 13% of breaches organizations experienced involved an AI model or application breach specifically, and 31% of those incidents caused operational disruption(s).⁶ Without operational maturity, deployments and related security concerns will continue to grow unchecked.

⁴ Bain & Company, "Technology Report 2025," www.bain.com/insights/topics/technology-report/

⁵ Cisco, "2025 Cisco Cybersecurity Readiness Index," <https://newsroom.cisco.com/c/r/newsroom/en/us/a/y2025/m05/cybersecurity-readiness-index-2025.html>

⁶ IBM Security and Ponemon Institute "Cost of a Data Breach Report 2025," www.ibm.com/reports/data-breach

The Three Roads

Governance, visibility, and control are three sides of the same problem (see Figure 3). Each road reinforces the others, and none of them works alone:

- Without visibility into where AI is running, what data it touches, and what actions it initiates, policy is aspirational. Visibility is the precondition for governance, not a parallel.
- Telemetry that surfaces an unauthorized agent or a misconfigured policy and produces no enforcement path is overhead without the effect. The investment in instrumentation only pays back when it is wired to action.
- Autonomous action based on stale or absent policy is exactly how the blast radius gets larger, not smaller. The faster the agent, the worse the damage when the rule it is enforcing no longer matches the organization's intent.

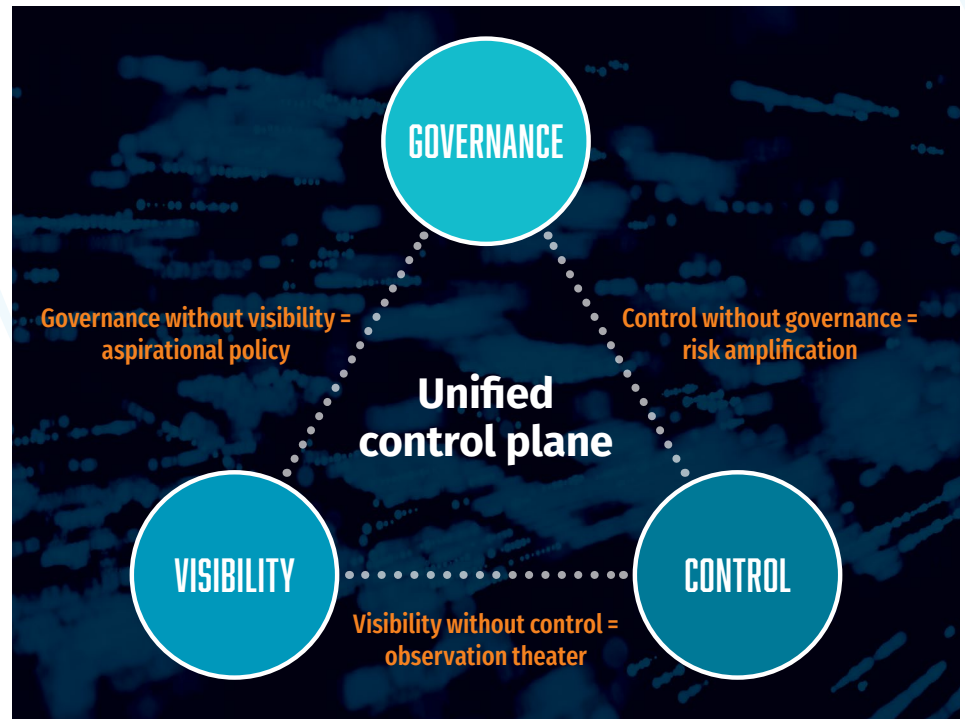


Figure 3. Governance, visibility, and control reinforce one another. Removing any one produces a characteristic failure mode.

The three roads run together but must work as a system to be effective.

Governance

The 2026 Regulatory Wall

The European Union's AI Act, Regulation 2024/1689, became binding for prohibited AI practices in February 2025 and for general-purpose AI obligations in August 2025. The next major step, which includes high-risk AI system obligations, takes effect on August 2, 2026.⁷ Penalties scale to 7% of global annual revenue, and the territorial reach mirrors the General Data Protection Regulation (GDPR): Any organization whose AI affects EU residents is in scope.⁸

For financial services, the Digital Operational Resilience Act (DORA), Regulation 2022/2554, has applied since January 17, 2025. This applies directly to AI systems used in financial operations:

- ICT risk management
- Incident reporting
- Resilience testing
- Third-party risk
- Information sharing

BaFin's December 2025 guidance on ICT risks in the use of AI made clear that AI systems must be embedded in DORA ICT frameworks rather than governed separately, and AI Act Article 9(10) expressly permits integrating AI risk management into existing DORA procedures. A timeline of European regulation is shown in Figure 4.

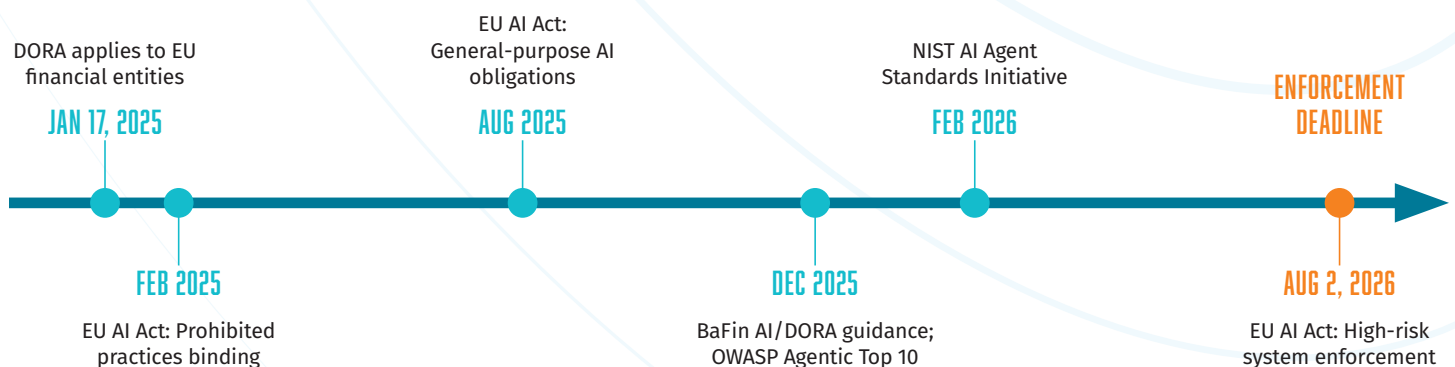


Figure 4. The Regulatory Timeline Converging on August 2, 2026, When EU AI Act High-Risk System Obligations Take Effect

⁷ European Commission, "AI Act," <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>

⁸ The European Commission's Digital Omnibus proposal of November 2025 could delay the high-risk obligations to December 2027, but as of this paper, that proposal is not yet law. For the purposes of this paper, we treat the August 2026 date as binding for planning purposes.

⁹ Federal Financial Supervisory Authority (BaFin), "Guidance on ICT Risks in the Use of AI at Financial Entities," December 2025, www.bafin.de

US sector frameworks layer on top of these:

- HIPAA's audit-trail requirements under 45 CFR § 164.312, GLBA, SOX, FedRAMP, and CCPA's deletion-on-request provisions all touch AI systems, whether they were written with AI in mind or not.
- Stanford Institute for Human-Centered AI's 2025 AI Index documented 59 new AI-related regulations issued by US federal agencies in 2024 alone, more than double the prior year.¹⁰

The regulatory surface is expanding faster than enterprise governance programs are scaling.

The Inventory Precondition

Every framework above assumes a defensible inventory of AI systems in use. Without inventory, classification is impossible. However, most enterprises do not have one. The IBM 2025 data showed that 20% of breaches involved shadow AI, with a \$670,000 cost premium per incident attributable to shadow AI specifically.¹¹

The Agent-Specific Governance Gap

Existing frameworks were written for AI systems that respond to inputs. Agentic AI systems initiate actions, choose tools at runtime, and accumulate state across interactions. Traditional policy breaks down when attempting to apply to agentic AI systems.

The *NIST AI Agent Standards Initiative*, announced in February 2026 as the first US federal framework specifically targeting autonomous AI systems, responds to a set of agent characteristics that distinguish them from earlier AI risk paradigms:

- Autonomous real-world actions with consequences
- Dynamic tool-switching at runtime that defeats statically evaluated policy
- Persistent memory as an attack surface that can be poisoned over time
- Non-deterministic behavior in which identical inputs may produce different actions¹²

Governance for agentic NetOps requires policy to be expressed in a form the agent can read and reason against at runtime, together with traceability from data through decision to action that supports audit reconstruction after the fact. Either the agent reasons against current policy or the policy is enforced reactively at a layer below the agent. In theory *and* practice, both are needed.

¹⁰ Stanford Institute for Human-Centered AI, "The 2025 AI Index Report—Policy and Governance," April 2025, <https://hai.stanford.edu/ai-index/2025-ai-index-report/policy-and-governance>

¹¹ IBM Security and Ponemon Institute, "Cost of a Data Breach Report 2025," July 2025, www.ibm.com/reports/data-breach

¹² National Institute of Standards and Technology, Center for AI Standards and Innovation, "Announcing the 'AI Agent Standards Initiative' for Interoperable and Secure Innovation," February 2026, www.nist.gov/news-events/news/2026/02/announcing-ai-agent-standards-initiative-interoperable-and-secure

Visibility

Visibility for AI is a two-sided problem: AI is a workload class that must be observed, and AI is the instrument that makes observation tractable at scale. That dual role is the most under-discussed dynamic in the market.

Blind Spots We Know About, But Cannot See

The data is all around us, and quite frankly, it's overwhelming:

- The *SANS 2025 Detection and Response Survey* identified the following: Endpoint-heavy detection produces network blind spots, particularly as workloads shift to cloud and mobile work expands the addressable perimeter.¹³
- The *2025 Verizon Data Breach Investigations Report*, drawn from 22,052 security incidents and 12,195 confirmed breaches, found that edge devices and VPNs now account for 22% of vulnerability exploitation targets—roughly an eightfold jump from 3% the prior year.¹⁴
- Mandiant's *M-Trends 2026* report, drawn from over 500,000 hours of frontline incident response, found that the time from initial access to handoff to a secondary threat actor has collapsed from more than eight hours in 2022 to 22 seconds today.¹⁵

AI as a Workload Class

Model-serving traffic, retrieval flows, and agent-to-agent (A2A) calls have traffic patterns that existing observability tooling was not designed to characterize. Bain's documentation of MCP server growth describes a new type of east-west traffic between agents and the tools they invoke. Vendors, including Kong, Microsoft, IBM, and Apache APISIX, have shipped MCP gateway products through 2026, signaling that this is an emerging operational discipline, not just research curiosity.

Without flow-level instrumentation, the basic questions about an AI-initiated call cannot be answered:

- Who initiated it?
- Under whose authorization?
- Against what data?
- With what scope?

Cisco's 2025 data found that 60% of organizations cannot see the specific prompts employees send to GenAI tools, and 41% lack mature controls over AI training data. Mandiant's *M-Trends 2026* documented AI-generated malware that rewrites itself hourly, only furthering the concept that signature-based defenses age out faster than they can be tuned.

¹³ SANS Institute, "2025 Detection and Response Survey," December 2025, www.sans.org/webcasts/sans-2025-detection-and-response-survey-webcast-forum

¹⁴ Verizon, "2025 Data Breach Investigations Report," April 2025, www.verizon.com/business/resources/reports/2025-dbir-data-breach-investigations-report.pdf

¹⁵ J. Kutscher, "M-Trends 2026: Data, Insights, and Strategies From the Frontlines," Mandiant Services/Google Cloud, April 2026, <https://cloud.google.com/blog/topics/threat-intelligence/m-trends-2026>

AI as the Visibility Instrument

The other half of the visibility story is where AI-assisted analysis is genuinely useful: telemetry volume reduction, cross-domain correlation, and anomaly framing. These are the SOC and NOC pain points documented in the *SANS 2025 SOC Survey*.¹⁶ Seventy-nine percent of SOCs operate around the clock, 69% still rely on manual reporting, and 40% use AI or ML tools (see Figure 5). However, satisfaction with those tools ranked *last* among all categories surveyed. The gap between deployment and usefulness is observably wide.

Organizations using AI and automation extensively in security operations saved an average of \$1.9 million per breach and reduced breach life cycle by 80 days vs. those that did not. The savings are real, but the honest limit needs to be named. AI visibility tools see only what they are fed. AI in the SOC inherits the same field of view used by the underlying tooling and instrumentation, unless expressly expanded. Garbage in, garbage out.

The staffing and tooling context



Figure 5. The Staffing and Tooling Context From the SANS 2025 SOC Survey. AI and ML Tools are used, but satisfaction with them ranked last among all categories surveyed.

¹⁶ C. Crowley, "SANS 2025 SOC Survey," SANS Institute, 2025, www.sans.org/white-papers/sans-2025-soc-survey

Control

Often, the goal is to augment AI, not replace it—reducing tool sprawl, responsibility sprawl, and the cognitive load that drives error. Control is where we prove this theory true.

Tool Sprawl as the Enemy of Control

The numbers tell the same story across a variety of sources:

- IBM data, cited by Verizon, shows the average enterprise has 83 security tools sourced from 29 vendors.¹⁷
- Ponemon's *Cyber Resilient Organization Study*—a 2020 baseline, but still the most-cited measurement in the category—put the average at 45 tools, with nearly 30% of organizations running 50 or more.¹⁸
- Saviynt counted 11 separate identity tools in the average enterprise, with 44% running multiple privileged access management products simultaneously.¹⁹

Cisco's 2025 data put hard numbers on the operational cost: 77% of organizations said complex security infrastructures with more than 10 solutions impede their ability to respond to threats.

Ponemon's survey presents perhaps the worst side of the results: Organizations running 50 or more tools were 8% less capable of detecting threats and 7% worse defensively than those running fewer.

**The data doesn't lie.
Survey after survey
finds that tool sprawl is
a *real* inhibiting factor
for security teams. More
tools, less security.**

¹⁷ R. Rehman, "Managing Cybersecurity Tool Sprawl," IBM, cited in Verizon, www.verizon.com/business/resources/articles/cybersecurity-tool-sprawl

¹⁸ Ponemon Institute, "Cyber Resilient Organization Study," 2020.

¹⁹ Saviynt, "Identity Sprawl Research," 2024, <https://saviynt.com/white-papers/identity-and-security-trends-and-predictions-2024-and-beyond>

Verified Automation

Operational guardrails for agentic NetOps need to be specified concretely. We suggest the following:

- Scoped authority defines the blast-radius limits within which an agent may act without human approval.
- Dry-run modes require the agent to propose and validate before executing, surfacing the planned action and its modeled effects for review.
- Tiered autonomy, stemming from TM Forum’s “in the loop,” “on the loop,” “out of the loop” model, sets which categories of work are autonomous (routine ticket triage), which require human ratification (network-wide changes), and which always escalate (novel or cross-domain incidents).²⁰
- Rollback paths make every action reversible within a defined commit window.
- Continuous policy validation requires the agent to check its intent against current policy state at decision time, not at code-deploy time.

The agentic AI characteristics named earlier, such as autonomous real-world action, dynamic tool-switching, persistent memory, and non-determinism, map directly onto why each of those guardrails is necessary rather than optional (see Figure 6). Each guardrail addresses a specific failure mode that the agent’s architecture makes possible.

Operational guardrail	Autonomous real-world action	Dynamic tool-switching	Persistent memory	Non-deterministic behavior
Scoped authority	✓			
Dry-run modes	✓		✓	✓
Tiered autonomy	✓			✓
Rollback paths	✓			✓
Continuous policy validation		✓	✓	

Figure 6. The Five Operational Guardrails for Agentic NetOps, Mapped to the Agentic Characteristics Identified by the NIST AI Agent Standards Initiative. Each guardrail exists because a specific architectural property of agents makes it necessary.

²⁰ S. Kurlkar, “The agentic era of telco AIOps: unifying IT and network operations with human oversight,” TM Forum, January 2026, <https://inform.tmforum.org/features-and-opinion/the-agentic-era-of-telco-aiops-unifying-it-and-network-operations-with-human-oversight>

Augmenting the Team

The 2025 SANS SOC Survey documents the staffing context and strengthens the argument for team augmentation, rather than replacement. Sixty-two percent of organizations say they are not doing enough to retain talent, and manual reporting eats hours that should be going to investigation. The case for AI assistance is recovering analyst time from work that machines do better.

In our experience and across our research, the areas where AI tooling presents benefits are limited, but highly effective: cross-tool correlation, first-pass triage, policy-diff review, change-impact analysis, and audit-evidence packaging. The common feature across these is that the cost of a marginal error is low and the error is recoverable. Where humans remain “in charge,” the deciders are the inverse: high-blast-radius changes, exception handling, novel threats—anything where the cost of being wrong is higher.

In this, we find the most useful heuristic for setting the human-to-AI handoff. It is more defensible in audit and incident review than abstract autonomy, because it forces the question: *If the agent is wrong, what does that cost, and how is it recovered?* Answering that question helps identify a concrete model for “who” should own the task.

True power comes when AI is designed to augment teams in areas where the cost of a marginal error is low and recoverable. Humans do best when the stakes are high, the threats are novel, and the cost of being wrong is asymmetric.

The 2026–2027 Roadmap

The name of the game, especially in this space, is consolidation. This means fewer tools, unified policy, common telemetry, and guarded orchestration across hybrid infrastructure. Roadmaps are being written to close the gap between adoption and operational maturity.

In the context of this paper, “unified” does not mean “one vendor for everything.” That premise is what produced the 83-tool average in the first place. Single-vendor platforms that solved part of the problem were then supplemented by point tools to cover the rest.

A better interpretation is coherent policy, telemetry, and orchestration layers across an entire tool stack. Heterogeneity is unavoidable in any large enterprise; thus, coherence is what gives an agent stability to reason against. Bain’s framing of agentic AI as (1) fit-for-purpose, (2) domain-specific, and (3) human-in-the-loop aligned.

A Practitioner Checklist

As we brought this paper together, our guiding concept was: Don’t just pose the problem; identify a solution or a path forward. Stating the problem is only, well, half of the problem. Thus, we present a checklist regarding AI-enhanced environments that CISOs and NetOps should address before the end of 2026:

- ☒ Where is AI deployed across the network?
- ☒ Which deployments fall under high-risk classifications?
- ☒ For every AI-initiated change, can a defensible audit chain be produced from data through decision to action?
- ☒ What is the blast-radius limit on any agent operating with autonomy, and how is it enforced rather than asserted?
- ☒ Is east–west traffic, including agent-to-agent and MCP, observed at flow level or only at the perimeter?
- ☒ What is the consolidation roadmap? Which tools are being decommissioned, and on what timeline? What coherent layer is replacing them?

Conclusion

Agentic NetOps is not a forecast. By the time most organizations get around to writing policy for autonomous systems, those systems are already running. They are opening tickets, proposing changes, and in a growing number of environments, executing them. The question facing most practitioners in 2026 is not whether to adopt agentic capability; that decision is largely already made. Rather, the question is whether the operational model *around* that capability is deliberate or accidental.

The three roads we identified in this paper only work as a system. Governance, visibility, and control must work together. In practice, the sequence matters less than coherence. You want an inventory you can defend, policy an agent can reason against at runtime, and telemetry that covers the east-west paths that agents use. Guardrails, such as scoped authority, dry runs, tiered autonomy, and continuous validation, convert autonomy from a “leap of faith” to an engineering discipline.

In any organization, the window for getting this right is measured in quarters. The EU AI Act’s high-risk obligations go into effect in August 2026, attacker handoff times are now measured in seconds, and agent-to-tool infrastructure is multiplying beneath operations teams faster than inventories can keep up. At no point does the clock pause to identify ownership. The teams that act quickly on implementation will meet the deadline as a milestone, rather than a crisis.

Sponsor

SANS would like to thank this paper’s sponsor:



About Our Author

AlgoSec Whitepaper Author

Matt Bromiley, Certified Instructor

Matt Bromiley is the Lead Solutions Engineer at LimaCharlie, where they build and maintain some of the best cybersecurity infrastructure capabilities. He has 12-plus years of experience in the cybersecurity industry making life difficult for cyber threat actors. He also serves as a GIAC Advisory Board member, a subject-matter expert for the SANS Security Awareness group, and a technical writer for the SANS Analyst Program. Matt brings his passion for incident response and leadership to the classroom as a SANS Instructor for **LDR553: Cyber Incident Management**.



 **READ MORE ABOUT MATT**

SANS | Research
Program

About the SANS Research Program

The SANS Research Program is a key initiative by the SANS Institute and a premier global provider of cybersecurity research and information. SANS Research Program is designed to provide cybersecurity practitioners and leaders with data-driven insights, thought leadership, and solutions that help them better understand and respond to evolving security challenges. All content is authored by SANS instructor experts from around the world who apply their years of experience from hands-on practitioner work in the field, advisory roles, and the classroom to provide education, guidance, and actionable insights that help make the cyber world a safer place.

To learn about sponsorship opportunities for research, content, and in-person or virtual events, email us at Sponsorships@sans.org or go to www.sans.org/sponsorship.