



2026 CLOUD SECURITY REPORT

Securing the AI Transformation

Why traditional network and cloud architectures are falling behind



Research by

Cybersecurity
INSIDERS

Executive Foreword

The speed of AI adoption is no longer the main story. The gap it has opened inside enterprise security is.

AI is not just introducing new tools; it is invalidating the assumptions security architectures were built on – how users behave, how applications communicate, and where control points exist. Security built on those assumptions is not failing because it is outdated, but because the environment it was designed for has fundamentally changed.

Employees now interact with external copilots from any device, applications exchange model-driven traffic at volumes that did not exist twelve months ago, and agents operate inside live systems with credentials, access, and autonomy. On the other side of that shift, attackers are weaponizing the same capabilities – faster phishing, AI-generated malware, adversarial inputs against models in production. Traditional security control models were not designed for this level of speed, autonomy, and distributed activity.

The research in this report puts a clear number on what most security leaders already sense. Strategy and investment are advancing, but architectural capacity is not keeping pace. In that gap sit the incidents, the shadow usage, the policy drift, and the growing blind spots in AI traffic.

At Check Point, secure AI adoption requires a prevention-first architecture, built closer to where AI interactions occur. This means consistent policy enforcement, shared visibility across environments, and real-time controls applied across prompts, data flows, and agent execution. Detection that occurs after a prompt has been processed, data has left the perimeter, or an agent has already acted is not protection – it is a log entry.

This gap is not inevitable. It will not be closed by adding more controls, but by building security into the architecture itself.



Roi Karo,
Chief Strategy Officer, Check Point

Overview

AI has moved past the experimentation phase and is now operating at enterprise scale. The survey shows two different adoption patterns happening at the same time. First, employee use of external Generative AI (GenAI) services is now a mainstream workplace reality. Second, organizations are increasingly building or embedding AI into their own applications, workflows, and agents.

The first challenge is governing access to external AI services such as ChatGPT, Copilot, and Gemini. These interactions create data exposure, policy enforcement, and visibility problems across browsers, endpoints, SaaS access paths, and remote work environments. The second challenge is securing enterprise-built or enterprise-operated AI systems. 70% of organizations run GenAI in production, and 64% have deployed AI agents inside live systems. But the security architecture surrounding those workloads has not kept pace. AI is being embedded into business workflows, applications, and infrastructure faster than enterprises can extend visibility, policy, and prevention across the environments where it operates. As AI moves deeper into applications and operations, security has to address not only what users send to AI, but also what AI systems can access, decide, and do.

The central finding in this report is the gap between strategic response and architectural readiness. 77% of organizations have changed their security strategy in response to AI, yet only 26% say they have the architecture to enforce it. That 51-point disconnect is already showing measurable consequences: 54% have reported an AI-related security incident, and closing it will require more than isolated point controls. It will require a unified security architecture that can carry policy, visibility, and prevention more consistently across users, applications, data flows, and hybrid environments.

Three patterns show how that gap is playing out.

- **Production ahead of control:** 70% already run GenAI in production, while 54% have experienced an AI-related security incident. AI is reaching live environments faster than enterprises are extending consistent control around it.
- **AI use is ahead of visibility and governance:** Only 5% of organizations report full visibility into AI tool usage across the organization. Most teams are trying to govern AI with an incomplete view of tools, agents, data flows, and runtime behavior.
- **Policy ahead of enforcement:** Only 14% say they have formal GenAI policies that are actively enforced and audited. At the same time, 42% say employees bypass controls when those controls slow them down, showing how quickly policy breaks down when it is not embedded into the normal path of work.

The pages that follow trace how this gap widens across infrastructure, access policy, application security, runtime enforcement, governance, and operating model design. The report closes with a maturity model and a five-step sequence for moving from fragmented AI security toward a unified hybrid mesh network security architecture with consistent policy, shared visibility, and distributed prevention across the hybrid environment.

AI Deployment Is Ahead of Defense

The pressure on security teams is not theoretical anymore. It is operational. On one side are employee interactions with external GenAI services and copilots. On the other are internally built or embedded AI capabilities that sit inside applications, workflows, and automated systems. 70% of organizations run GenAI workloads in production, and 83% say securing them is harder than securing traditional applications. That gap widens as AI moves from generating content to taking action. AI agents have reached pilot or production in 64% of enterprises, and 12% have granted them privileged access to core systems. The security issue has shifted from what employees ask AI to do toward what AI systems are allowed to do on their own.

That change matters because most security architectures were built for a different model: human-driven access, known SaaS patterns, and more predictable application behavior. They were not designed for AI interactions that are dynamic, API-mediated, and increasingly autonomous. Security teams have to govern both types of exposure at once: employee access to external AI tools and machine-driven actions inside enterprise workflows.

AI Is Already in Production

► How would you describe your organization's current adoption of Generative AI?

Core to business operations

10%

Broad production

25%

Limited production

35%

Experimentation or pilot

24%

No active use

6%

70%

organizations are running GenAI in production



83%



say securing GenAI is harder than traditional applications

64%

of organizations have AI agents in pilot or production, including 12% with privileged access to core systems

Enforcement maturity tells the same story: only 14% have AI security policies that are both enforced and audited. Organizations moving faster are treating agent deployment as an architectural control decision: mapping credentials, restricting access, and defining which actions require human approval before execution. Agent adoption is scaling faster than most enterprise control models were built to contain.

The Threat Surface Is Already Live

AI adoption at production scale creates exposure at production scale. 54% of organizations have confirmed at least one AI-related security incident, while another 24% suspect an incident but lack the telemetry to confirm it. Together, that means 78% have either experienced confirmed AI-related security impact or cannot rule it out. That uncertainty is evidence in itself: many detection environments are not tuned for AI-specific threats.

The reported incident types reinforce that this is not one problem. When respondents were asked what kinds of AI-related incidents they had experienced, the most common answers were unauthorized or shadow AI usage discovered (41%), AI-generated content used in an attack such as phishing or deepfakes (37%), and sensitive data leaked to or through AI services (32%). That spans both sides of the AI security problem: governing employee use of external AI services and protecting enterprise environments from AI-enabled attacks or weaknesses in enterprise-operated AI systems.

That mix is what makes the threat surface difficult to govern. AI-related activity can resemble legitimate AI traffic. API calls, model queries, and outbound requests to AI services may all appear routine at the network layer unless the inspection point can evaluate what the interaction is doing. Recent threat research has already shown that attackers can disguise malicious activity inside AI-like API traffic, making shallow inspection increasingly unreliable in high-volume AI environments.

Confirmed and Suspected AI Incidents

► What is the most significant GenAI-related security incident your organization has experienced?

78% of organizations reporting confirmed or suspected AI security incidents



54% of organizations have confirmed at least one AI-related security incident

That blind spot will widen as AI traffic becomes more common across enterprise environments. The more normal outbound AI traffic becomes, the easier it is for malicious activity to blend into it. The practical response is to extend inspection deeper into AI traffic paths and treat outbound calls to LLM services as sensitive transaction channels rather than routine web activity.

The 51-Point AI Readiness Gap

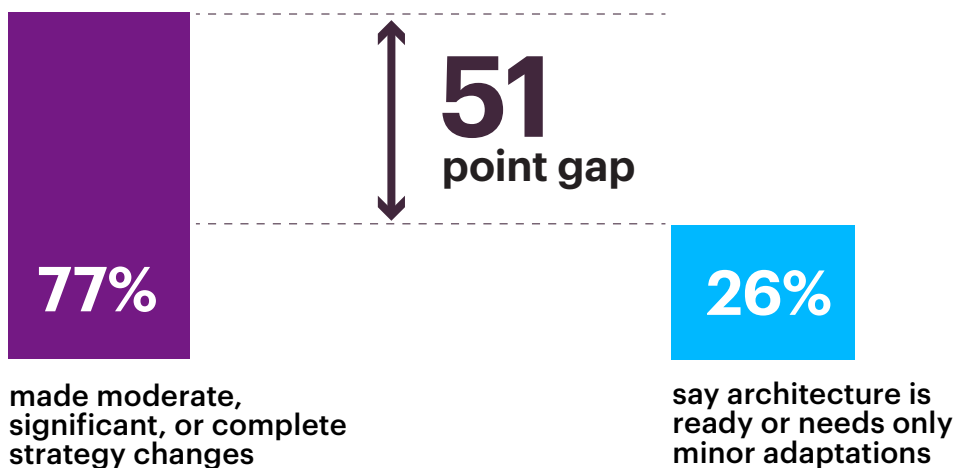
The incidents confirm the urgency, but the deeper problem is architectural. 77% of organizations have changed their security strategy in response to AI: rewriting acceptable use policies, standing up AI governance boards, and redirecting budget toward AI-specific controls. Yet only 26% say they have the infrastructure to execute that shift. Strategy is intent; architecture is capacity. In AI environments, that capacity depends on whether policy, visibility, and prevention can operate as one system.

That 51-point gap between strategic commitment and architectural readiness is the widest finding in the survey, and it explains why the same pattern repeats across the rest of the data. AI workloads are crossing cloud, datacenter, and SaaS environments where firewall rules, DLP policies, and access controls often change at the boundary, or disappear entirely. Security may exist in one domain, but the architecture does not carry it consistently into the next.

Strategy Has Moved Faster Than Architecture

► How significantly has GenAI adoption changed your organization's overall security strategy?

► Can your current security architecture support AI-driven workloads without significant redesign?



Organizations closing this gap are rebuilding strategy and enforcement in parallel, starting with a unified policy architecture that follows AI workloads across environments instead of relying on each domain to enforce its own rules. In practice, that points toward a hybrid mesh network security architecture: policy defined once, distributed consistently across cloud, datacenter, SaaS, and user access paths, with prevention enforced close to where workloads, traffic, and user interactions occur. The organizations moving fastest are treating policy and enforcement as one coordinated operating model, not as separate controls managed environment by environment.

The AI Blind Spot

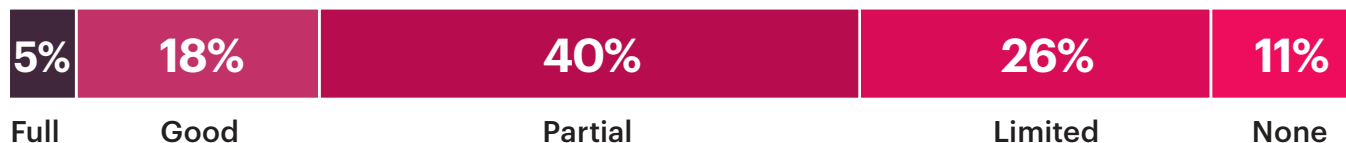
The AI readiness gap starts with a more basic failure: most security teams still cannot fully see what they are trying to protect. Only 5% have full visibility into which AI tools employees are using, what data those tools are accessing, and where that data goes once it enters an AI workflow. That leaves 95% making policy and enforcement decisions from an incomplete map of their AI environment.

That blind spot shows up in more than discovery. A separate finding shows that only 5% said their security tools can reliably distinguish legitimate AI activity from suspicious or unauthorized usage. The visibility problem is therefore twofold: security teams often do not know which AI tools are being used, and they often cannot determine whether the activity they do see is legitimate or risky.

Traditional discovery tools were built to find known applications, managed endpoints, and cataloged SaaS usage. AI activity often bypasses those assumptions. A browser-based assistant may leave little endpoint evidence. An LLM API call from a script or notebook may never appear in SaaS discovery. An agent using an existing service account can look like normal system activity. Without AI-specific telemetry, much of this activity remains difficult for the SOC to classify correctly.

Only 5% have full visibility into AI tools and services

► How much visibility does your security team have into all GenAI tools and services being used across your organization?



Only 5% can reliably distinguish legitimate AI activity from suspicious usage

Organizations closing this visibility gap are building AI-specific inventories before they attempt to enforce broad policy: separately mapping employee access to external AI tools, enterprise AI applications and APIs, non-human identities used by agents, and the sensitive data flows that connect them.

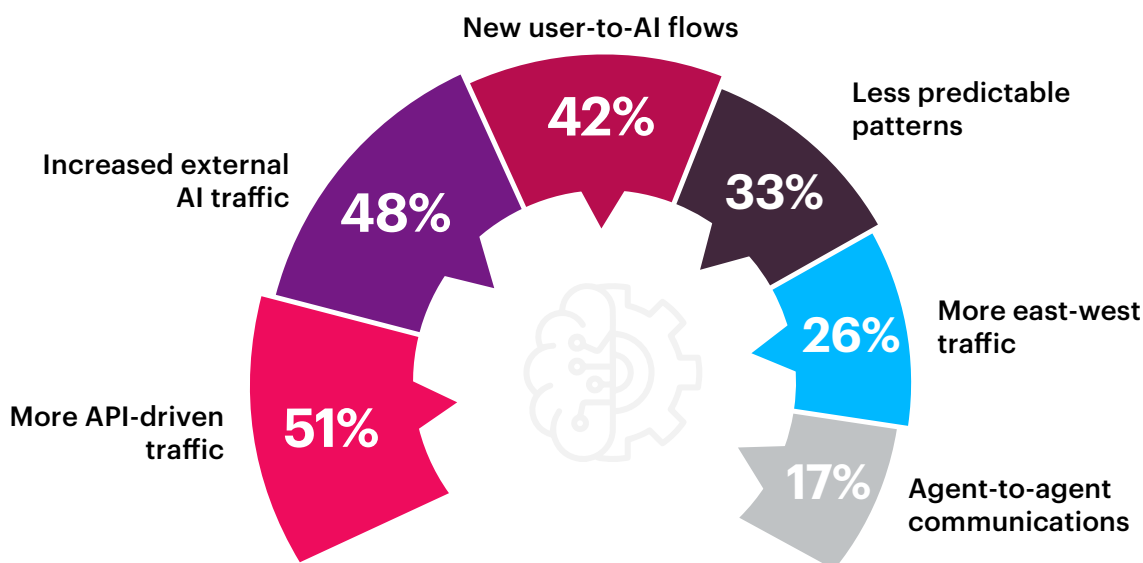
Infrastructure Security Was Built for a Different Traffic Model

AI is changing traffic patterns faster than most infrastructure security stacks were designed to handle. In the survey, 51% of organizations reported more API-driven traffic, 48% reported increased traffic to external AI services, 42% reported new user-to-AI communication flows, 33% reported less predictable traffic patterns, and 26% reported more east-west traffic. These are not edge cases. They are signs of a different operating model.

That misalignment shows up in control coverage. Only 24% said their current network security tools can fully inspect AI traffic without degrading application performance, leaving 76% with inspection gaps, performance tradeoffs, or limited confidence in their current controls. At the same time, 67% of organizations report fragmented security policies across their hybrid environments, and 64% said their architecture needs moderate or significant redesign to support AI workloads.

How AI Has Changed Network Traffic Patterns

► How has GenAI affected network traffic patterns in your environment?



76% face AI traffic inspection gaps or performance tradeoffs



Only 24% can fully inspect AI traffic without performance impact

That is why AI pressure shows up as an infrastructure problem before it becomes a product problem. Security architectures built around predictable user sessions and stable application paths are now being asked to govern API-heavy, service-mediated, and sometimes agent-driven traffic across multiple environments. The organizations adapting fastest are redesigning around the new traffic model: shared policy, deeper inspection where it matters, and enforcement that can span external AI access, internal application flows, and east-west movement without collapsing under performance strain.

AI Is Reshaping Datacenter Security

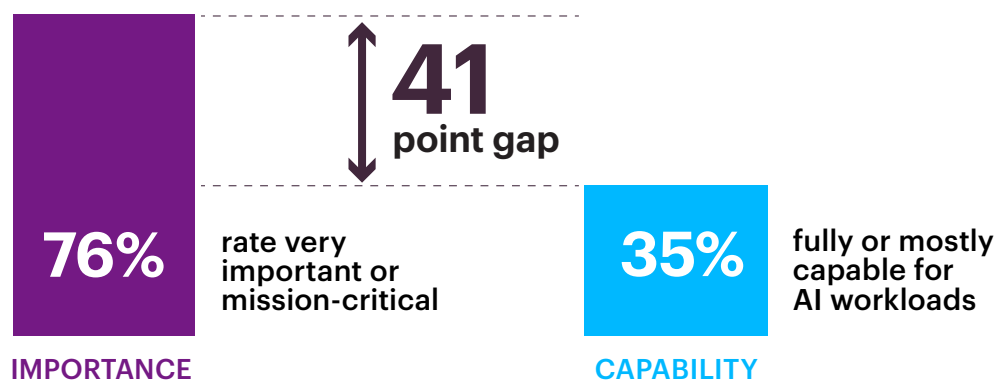
Fragmentation across cloud and SaaS is only part of the infrastructure challenge. AI workloads are also expanding into on-premises and private environments. 29% of organizations are already migrating AI training and inference to their own datacenters, with another 49% planned or under consideration. As AI becomes more data-intensive and operationally important, more organizations are placing compute closer to the systems, data, and regulatory boundaries it depends on.

That shift aligns with the broader deployment picture: 52% said their AI models and AI-powered applications primarily run across a hybrid mix of cloud, on-premises, and third-party services. These environments place different demands on security than traditional enterprise infrastructure did. AI training and inference create high-volume traffic flows, more east-west communication, and tighter coupling between storage, orchestration, model serving, and downstream applications. Not surprisingly, 76% rate datacenter perimeter security as critical for AI workloads, but only 35% say their current perimeter can actually handle them. The gap is now showing up in the architecture itself.

Datacenter Perimeter: Importance vs. Capability

► How critical is datacenter perimeter security for protecting your AI training and inference infrastructure?

► How would you rate your organization's current ability to secure the datacenter perimeter specifically for AI training and inference workloads?



29% already moving AI workloads to on-premises/private cloud

49% planned or under consideration

As AI workloads become more hybrid, the datacenter starts looking less like a separate zone and more like one enforcement domain inside a broader hybrid network security architecture. Controls have to extend across north-south and east-west traffic, with inspection and prevention placed close to the workloads and traffic paths they govern, without becoming a throughput bottleneck.

Inconsistent Governance of Employee AI Access

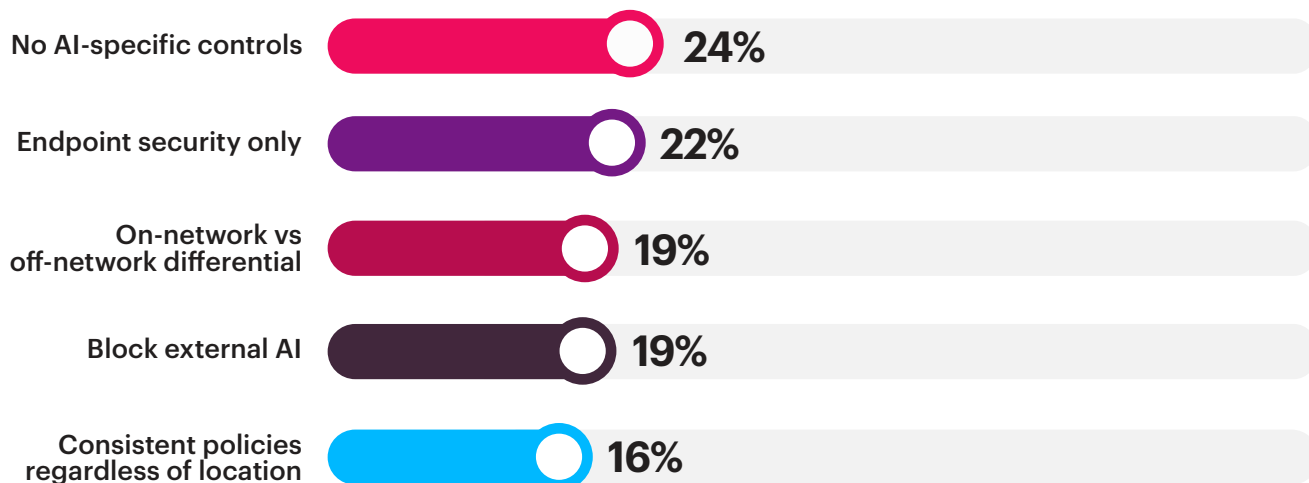
The infrastructure layer governs traffic between environments. The workspace layer governs how employees actually interact with AI day to day. Right now, that control model remains highly inconsistent.

Organizations have not converged on a single access model: 24% have no AI-specific access controls, 22% still rely primarily on endpoint agents, 19% apply different rules on-network and off-network, 19% block external AI tools outright. Only 16% enforce the same controls consistently. Five strategies, none dominant.

The result is that the same employee can encounter very different AI security depending on where they are working and how they connect.

Five Access Strategies, No Dominant Model

► How does your organization secure employee access to GenAI tools and services (e.g., ChatGPT, Copilot, internal AI)?



The deeper lesson is architectural. AI access policy cannot depend primarily on network location if the same tools, data, and users move freely across office, remote, and SaaS environments. Rather, it requires policy that can identify AI destinations, distinguish sanctioned from unsanctioned services, restrict use of personal tenants where needed, apply data-handling rules to uploads and downloads, and incorporate user and device context.

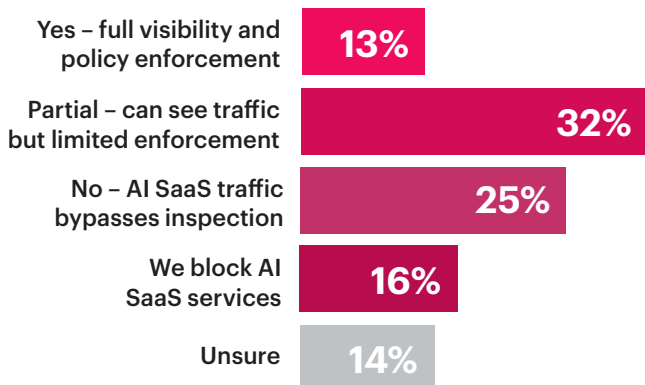
Endpoint and SaaS Coverage Falls Short

Fragmented policy is only part of the problem. The other is control coverage. Only 13% can fully enforce security policy on traffic to AI SaaS services like ChatGPT, Copilot, and Gemini. The rest operate with gaps, partial coverage, or traffic that bypasses inspection entirely.

39% say their endpoint security tools do not cover AI applications. Detection of shadow AI on corporate devices is similarly weak: only 12% report real-time detection and alerting, while 38% rely on periodic scanning or audits and 34% rely on policy compliance alone.

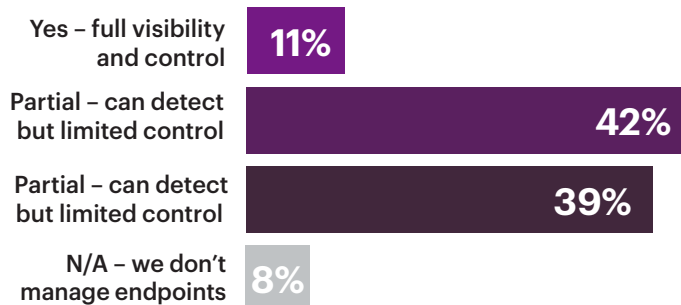
AI SaaS Traffic Enforcement

► Can your security controls inspect and enforce policy on traffic to AI SaaS services (e.g., OpenAI, Anthropic, Google AI) in real-time?



Endpoint and Browser AI Control

► Can your endpoint security tools detect and control unauthorized AI applications or browser-based AI tools on managed devices?



AI coverage is fragmented across both planes:

Only 13% can fully inspect and enforce AI SaaS traffic, and only 11% can fully detect and control endpoint/browser AI usage

This is where coverage gaps compound each other. Browser sessions, direct API calls, file uploads, downloads, copy-and-paste behavior, and personal-account sign-ins do not all show up at the same control point. Governing external AI services therefore depends on combining several capabilities that many organizations still operate separately: access policy for AI destinations, application-level policy for sanctioned versus unsanctioned services, tenant-aware controls, DLP for sensitive uploads and downloads, security event telemetry, and device or posture context. Without that combined coverage, employee AI usage remains visible in fragments and enforceable only in fragments.

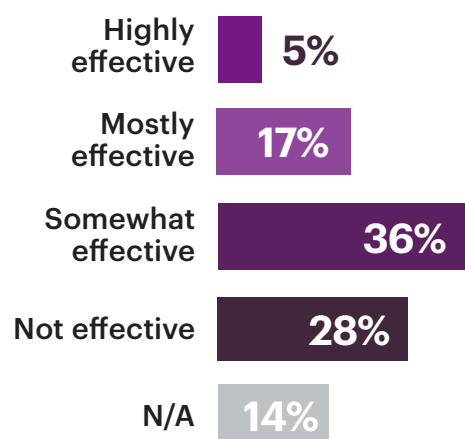
AI Traffic Is Exposing WAF and WAAP Limits

Access controls govern how employees reach AI tools, but at the application layer, the challenge changes to protecting enterprise AI applications, APIs, and agents. This is where models process prompts, generate outputs, and interact with sensitive data and services. The controls are often present, but AI traffic is exposing how poorly many of them fit this new interaction model. Only 22% rate their current WAF or WAAP tools as effective for GenAI-specific attacks such as prompt injection. A much larger share reports limitations: 64% say their tools are somewhat to not effective at all, while 28% say their current tools were not built for AI threats. The operational impact is already visible, with 71% reporting increased WAF false positives since GenAI adoption.

Coverage of AI assets is also uneven: 38% said their application security controls protect public-facing AI web applications, 32% protect APIs connecting to AI services, 28% protect LLM endpoints, and 20% protect AI agents, while 21% reported no AI-specific application security at all.

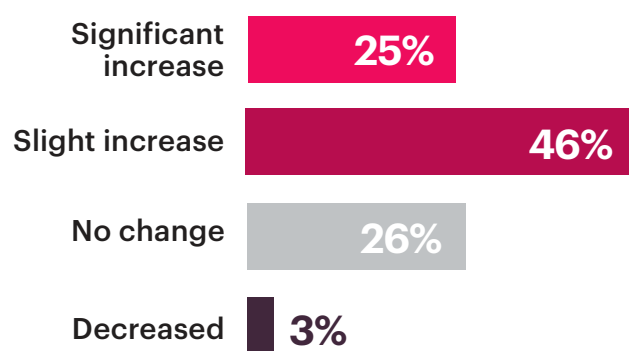
WAF/WAAP Effectiveness for Detecting AI-specific Attacks

► How effective are your current WAF/WAAP tools at detecting GenAI-specific attacks like prompt injection?



WAF False Positive Trends Since AI Adoption

► Has GenAI increased false positives in your WAF?



Only 22% rate WAF/WAAP tools effective against GenAI-specific attacks
71% report increased false positives since GenAI adoption

That mismatch is structural. Traditional WAF logic was tuned for human-driven web traffic, known attack signatures, and more predictable request patterns. GenAI applications introduce long prompts, streaming responses, model-specific APIs, and service-to-service interactions that do not map neatly to those assumptions. Organizations closing this gap are expanding application security beyond classic browser traffic: updating inspection logic for AI payloads, extending control into API and model-serving paths, and connecting application-layer enforcement back to the broader policy, identity, and data-protection model.

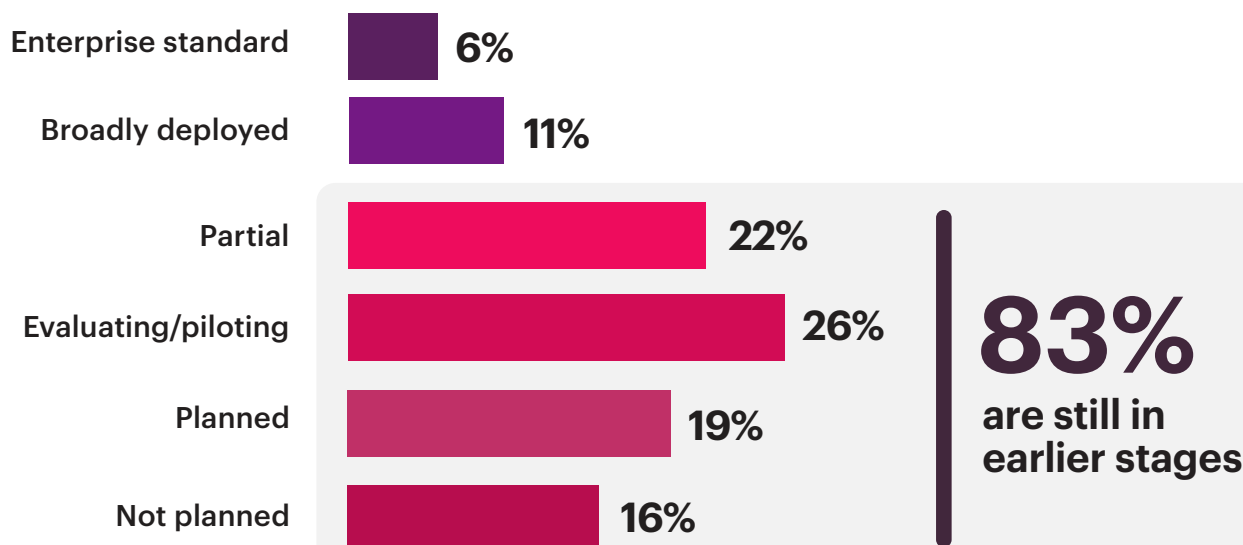
Runtime Controls Still Lag Behind Deployment

WAF degradation affects traffic at the application edge. Inside the application, the exposure is more direct. Models are reading sensitive data, generating outputs, and invoking tools and APIs in real time. Without runtime guardrails, those interactions remain weakly governed at the point where decisions are actually made. Only 17% have broadly deployed runtime LLM controls such as input validation, output filtering, and tool-use authorization across their applications.

The testing picture is not much stronger. 56% have no formal security testing process for GenAI applications, or test only ad hoc. That means many AI applications are reaching production without structured prompt-injection testing, without repeatable adversarial evaluation in the CI/CD pipeline, and without consistent retesting after models or workflows change. The result is predictable: organizations are deploying AI functionality faster than they are validating whether it can be safely governed in production.

Runtime LLM Control Deployment Maturity

► Does your organization have runtime controls that inspect and enforce policy on LLM inputs and outputs in real-time?



56% have no formal testing process for GenAI applications, or test only on an ad-hoc basis

The priority now is to treat runtime control as a production requirement, not a later enhancement. Input validation, output filtering, tool authorization, and repeatable security testing need to be in place before AI applications scale into business-critical workflows. Once models operate in production, runtime enforcement becomes the control layer that determines whether the application can be governed safely at all.

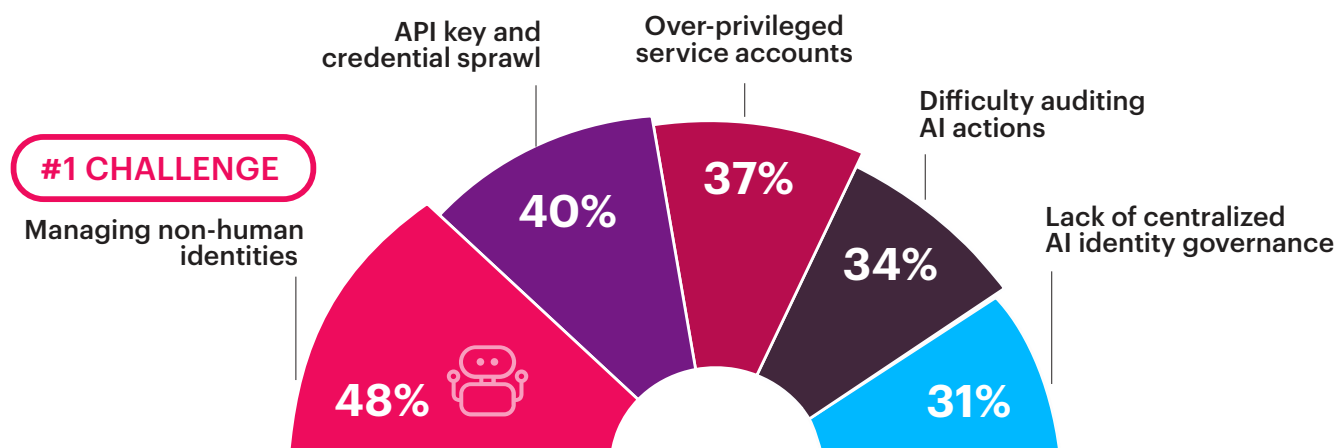
The New AI Attack Surface

AI is not only stressing existing controls, it is also creating access paths that many enterprises were not built to govern well. Non-human identity is the clearest example. 48% of organizations rank it as their leading AI-related identity challenge. As AI agents begin authenticating with service accounts, API keys, and delegated permissions, enterprises are having to govern machine-driven access patterns that do not resemble human behavior or fit neatly into human-centric IAM models.

That exposure extends beyond identity into the AI supply chain. 46% of organizations have no structured security assessment process for AI vendors, and only 7% scan AI models for vulnerabilities or malicious code before deployment. Many enterprises are deploying models they have not inspected, from vendors they have not assessed, into applications they have not fully tested. Recent package-compromise incidents in the AI ecosystem have shown how easily credentials, tokens, and downstream access can be exposed when trust is extended too early across the software and model supply chain.

Top AI Identity and Access Challenges

► Which AI-related identity and access challenges does your organization face?



46% have no structured security assessment for AI vendors

Only 7%
always scan AI models
before deployment

The lesson is that AI expands trust relationships faster than most governance models are adapting. Organizations addressing this exposure are moving non-human identity governance, credential discipline, vendor review, and model assurance into the core security program now, before these dependencies harden into default enterprise behavior.

AI Is Creating a New Data Exposure Path

Data movement is the broadest AI security exposure because it cuts across all major AI use cases in this report. Every time an employee pastes text into ChatGPT, data leaves through an external service boundary. When an internal AI application retrieves sensitive records, sends them to a model, and returns generated output downstream, the movement is different, but the governance question is the same: what data moved, where did it go, and what policy was applied at each step?

This survey shows how unsettled the policy layer still is. 25% of organizations permit source code in AI tools today, exposing proprietary business logic, internal configurations, and competitive IP inside external systems with limited exfiltration control. 44% cannot trace where sensitive data goes once it enters an AI workflow: no clear lineage from source system to model interaction to output. Only 15% have deployed DLP controls specifically configured for AI data flows. Many enterprises are allowing sensitive information into AI systems without the lineage, inspection, and enforcement needed to govern what happens next.

Data Lineage Visibility Through AI Processing

- Can your organization track the flow of sensitive data from source systems, through AI processing, to output?

61% of organizations reporting no visibility or are unsure



Only 15% have AI-specific DLP deployed and enforced

25% permit source code in GenAI tools



AI security becomes data-centric at this point. Source code, internal documents, structured records, and prompt content all need to be governed as high-value data flows from the moment they enter an AI interaction. AI-specific DLP, lineage tracking, and policy enforcement at the point of submission have to be in place before enterprises can claim meaningful control over how sensitive information moves through AI systems. Without those controls, security teams are left to infer data movement after the fact instead of governing it at the moment it happens.

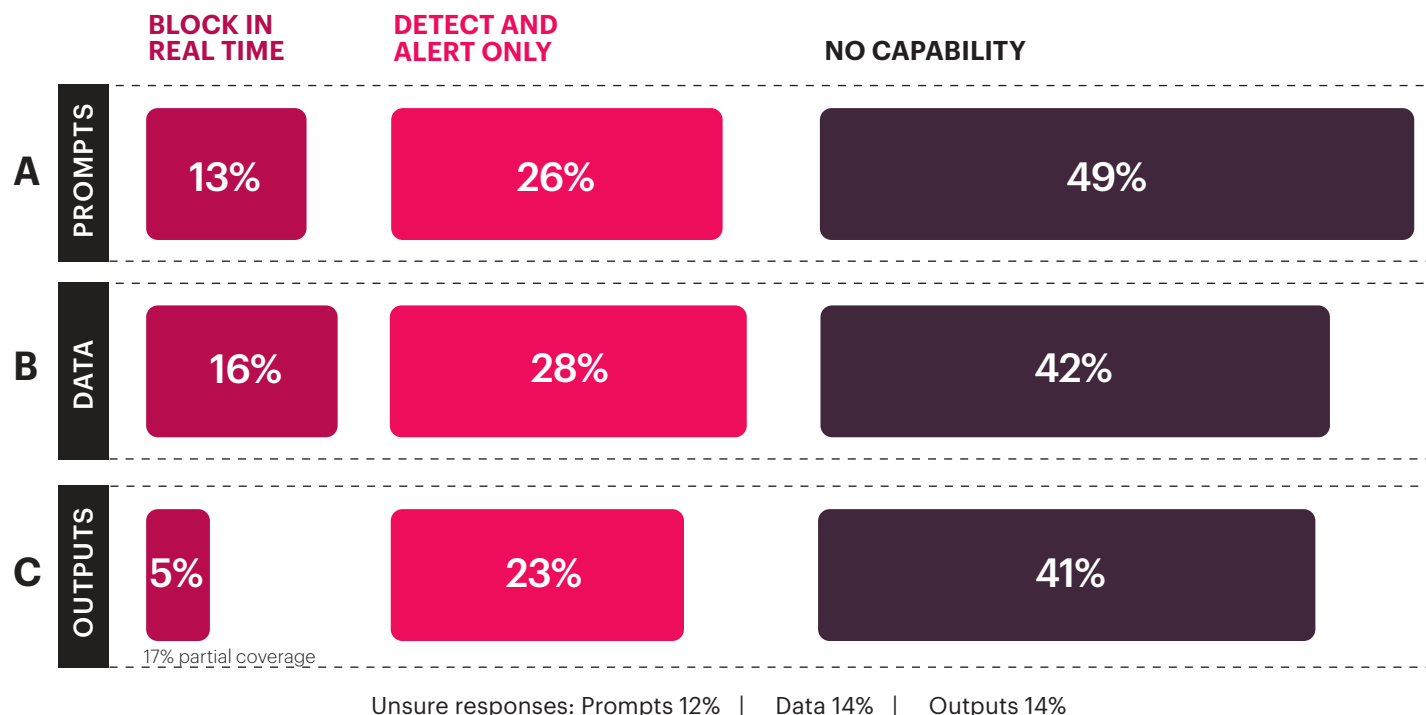
Detection Without Prevention

Every attack surface examined so far shares the same structural weakness: security tools can often detect AI-related threats, but very few can block them in real time. Alerts fire after the data has left, the output has been generated, or the access has already been used. Detection without prevention produces evidence after the event, not control at the moment it matters.

The gap is measurable across all three enforcement points. Prompt security is the most mature, and even there only 13% can block a malicious prompt or jailbreak attempt before it reaches the model, while 26% can detect and alert, but not block. 16% can block sensitive data from reaching AI services, while 28% detect and alert without enforcement. Outputs are the weakest point of all: only 5% can reliably block unsafe AI-generated content before it reaches users. Across prompts, data flows, and outputs, the pattern is the same. More organizations can observe the risk, but cannot stop it.

Prevention Gap Across Prompts, Data, and Outputs

- ▶ **A) Can your security tools automatically detect and block malicious or manipulative prompts sent to AI services in real-time?**
- ▶ **B) Can your security tools prevent sensitive or regulated data from being sent to AI services in real-time?**
- ▶ **C) Can your security controls block or filter unsafe AI-generated outputs before they reach users or downstream systems?**



This is where AI security architecture has to change. A prevention-first model needs to sit in the data path, not alongside it as a monitoring layer. The first priority should be data flows, because they represent the broadest and highest-volume exposure surface; from there, the same enforcement model can extend to prompt filtering and output validation. Prevention becomes real when inspection engines are positioned to act at inference speed, not simply record what already happened.

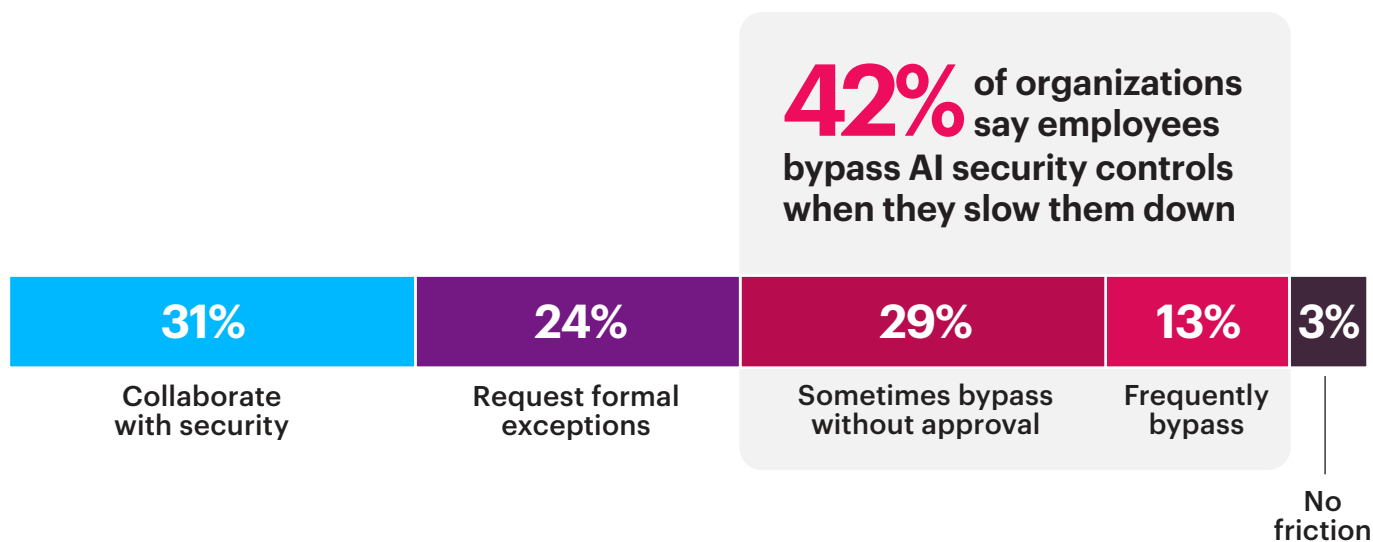
Employees Are Working Around AI Security

The earlier gaps were architectural. This one shows how those architectural weaknesses reflect in user behavior. 42% of organizations say employees bypass AI security controls when those controls slow them down. They paste sensitive data into a personal AI account, use a browser-based tool the endpoint agent does not cover, or take whatever path feels faster than the approved one. This behavior signals something structural: governance is sitting too far above the workflow to shape what users actually do.

Security teams are aware of the tension. 51% believe they are enabling AI adoption, yet 21% say slowing AI adoption for security reasons has already cost them competitive edge. That gap matters. It means policy may exist, executive sponsorship may exist, and governance committees may exist, while employees still choose faster paths outside the approved control plane. In practice, friction is converting governance intent into operational drift.

How Teams Respond When Security Creates Friction

► When security controls create friction or latency for AI workloads, how do teams typically respond?



Governance starts working when enforcement is embedded into the normal path of work rather than layered on top of it. Approved AI access needs to be easier to use, consistently governed across environments, and tied to the same policy logic whether the interaction happens through browser, API, or managed device. When secure workflows are usable by default, user bypass behavior starts shrinking as a control problem.

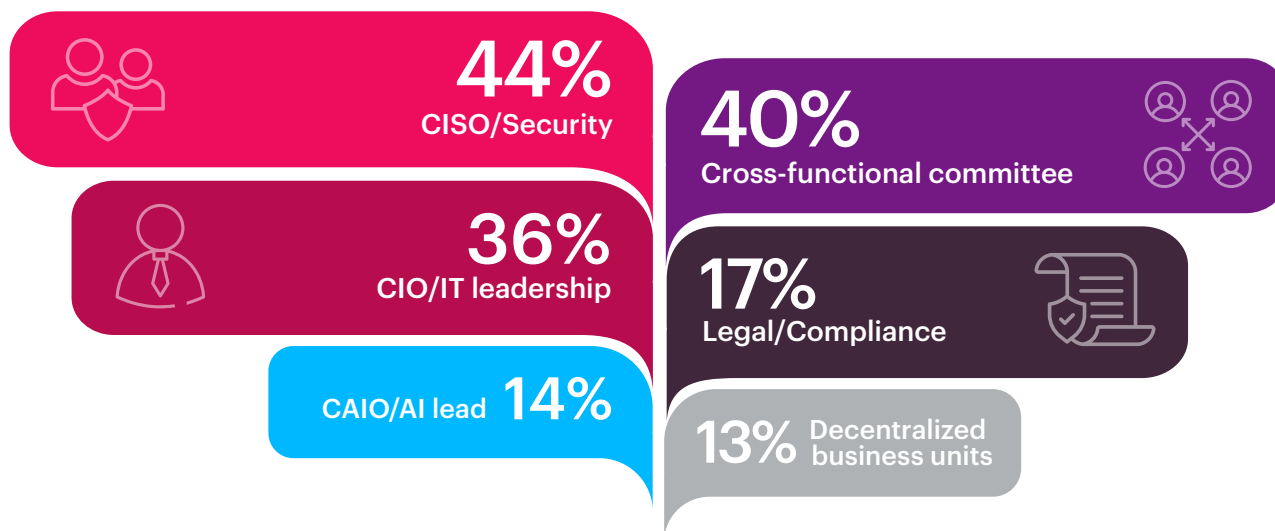
Accountability Without Authority

Who owns AI security? In many organizations, the answer is still too broad to be operational. Accountability is spread across the CISO (44%), cross-functional committees (40%), and CIO or IT leadership (36%). Only 14% have a dedicated AI security leader, and only 14% have AI security policies that are both enforced and audited. Together, those findings point to the same operating problem: accountability is often present in principle but weak in execution.

The harder question is whether anyone can turn policy into consistent control. 45% of organizations have documented AI security policies. At the same time, regulatory pressure is increasing via AI-specific requirements and governance frameworks such as the EU AI Act and NIST AI RMF, alongside sector-specific obligations. Only 7% say they are fully prepared, while another 29% are aware of requirements but have not operationalized compliance. That leaves many teams trying to answer governance and audit demands without a consistent enforcement model underneath them.

Who Owns AI Security Risk

► Who holds the primary accountability for AI Security risks and policy enforcement in your organization?



Governance execution gap: 45% have documented AI security policies > 14% actively enforce and audit them

Accountability starts working when policy is attached to operational authority and distributed through the same control architecture that governs the environment day to day. One owner does not need to make every decision, but one function does need authority to define policy once, distribute it consistently across cloud, datacenter, endpoint, SaaS, and AI control layers, and produce evidence that enforcement is actually happening. Governance becomes durable when compliance is generated by the architecture itself, rather than reconstructed manually after the fact.

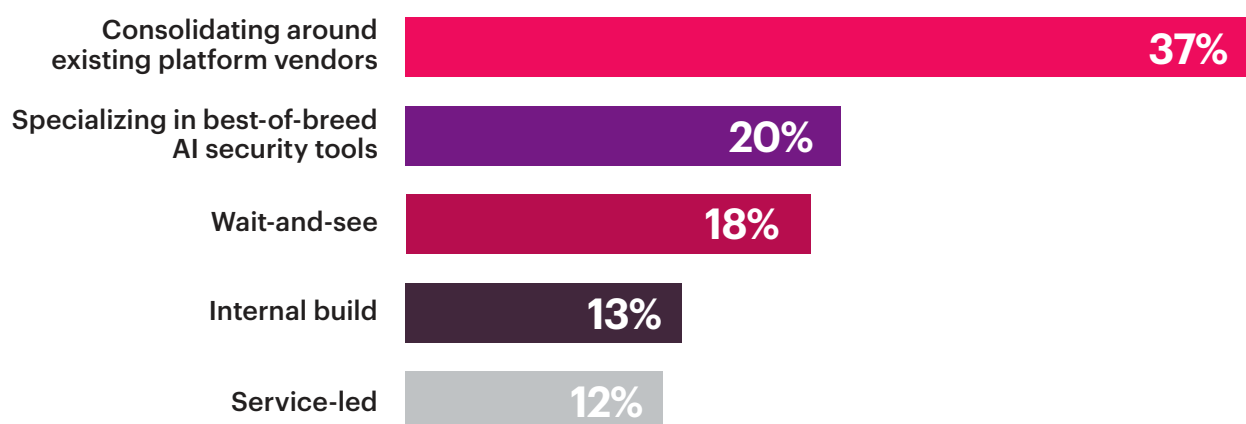
The Market Is Moving Toward a Broader Operating Model

The response to AI security is shifting from isolated fixes toward broader operating models. 75% of organizations have changed their underlying security architecture strategy in response to AI, and 52% are increasing dedicated AI security budgets. The direction of that spend matters: the market is consolidating toward fewer, broader control planes and unified operating models that can govern AI activity consistently across the hybrid environment.

That direction is clear in the investment data. 86% rate unified security management across datacenter, cloud, and edge as critical for AI workloads. 37% are actively consolidating, nearly double the 20% still investing in best-of-breed point solutions. In operational terms, that means unified management, fewer policy consoles, more consistent enforcement outcomes, and less dependence on disconnected security silos. For many organizations, the practical path forward is not simply adding another AI-specific control. It is consolidating toward integrated platforms that can manage policy centrally, enforce it consistently, and reduce friction for users and security teams alike.

Primary AI Investment Direction Over the Next 12 Months

► Which statement best describes your organization's primary investment direction for AI security over the next 12 months?



86% rate unified security management across datacenter, cloud, and edge as very important or critical for AI workloads

Investment priorities reinforce the same lesson. Shadow AI discovery leads at 31%, ahead of AI-powered SOC automation at 25% and adversarial defense capabilities at 19%. That sequencing suggests many organizations are trying first to regain control over what is already happening in the environment. As they do, the operational pressure favors fewer disconnected control planes and more shared policy and management layers.

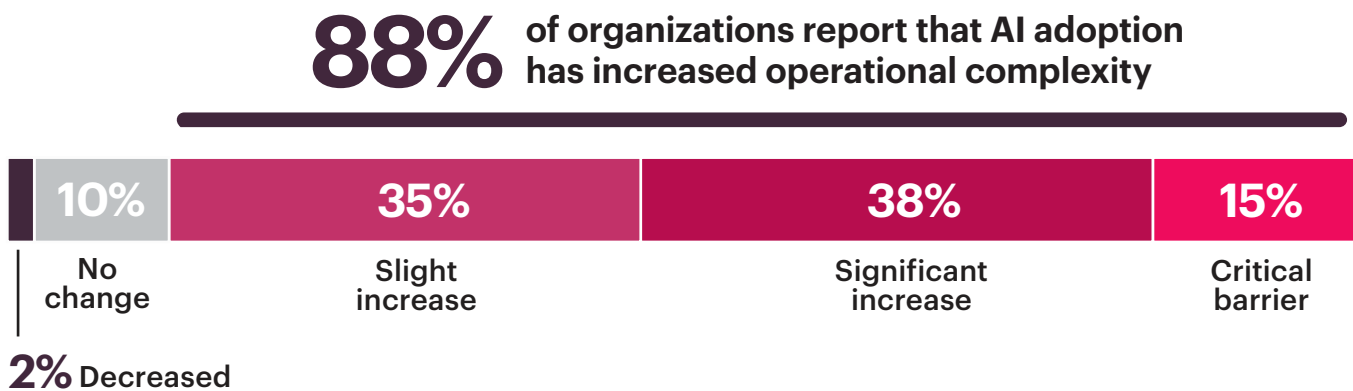
The Aspiration-Execution Gap

The market is moving in the right direction, but complexity is still moving faster. 88% of organizations report that AI adoption has increased the operational complexity of their security program. The leading barriers are the same weaknesses that have appeared throughout this report: lack of visibility at 28%, skills gaps at 25%, and disconnected tools at 16%. AI is exposing and amplifying these weaknesses as more users, applications, agents, data flows, and security workflows become AI-enabled.

Organizations increasingly recognize the need for unified management, but execution remains difficult, and it's visible in the numbers. 86% of organizations say unified security management is critical for AI workloads. Yet only 37% are actively building toward it. That leaves a large middle ground where the target architecture is understood, but the operating model has not caught up. AI is increasing the urgency of consolidation faster than most enterprises are reducing the fragmentation beneath it.

AI's Impact on Operational Security Complexity

► How has GenAI adoption affected the operational complexity of your security program?



Execution now has to catch up with intent. The next step is not simply to add more AI features into the stack. It is to reduce policy fragmentation, clarify ownership, and connect visibility, prevention, and response across the control points that AI has exposed. Otherwise, AI improves isolated parts of the security program while making the overall operating model harder to manage.

The AI Security Maturity Profile

Every gap in this survey points back to the same question: how mature is your organization's AI security, and what kind of operating model is needed to close the gaps? This matrix maps the answer across six capability areas. At the low end, AI security remains fragmented and reactive. At the high end, it begins to resemble a hybrid mesh network security architecture with shared visibility, a single policy framework, and distributed enforcement across the hybrid environment. Find the closest description that matches your current operating reality. The column to its right shows the next level of maturity.

CAPABILITY	REACTIVE	MANAGED	ADAPTIVE
Governance & Risk Alignment	Policies documented but not enforced. Accountability diffused across multiple owners. Controls bypassed when they create friction.	Accountable owner defines policy. Governance requirements are translated into enforceable controls across key environments.	Policy defined centrally and applied automatically as new AI workloads, agents, and data flows are deployed.
Visibility & Situational Awareness	AI tools, agents, and data flows largely invisible. No unified asset inventory.	Unified AI telemetry across all environments from a single observability plane. Shadow AI discoverable.	New AI workloads discovered and classified automatically as they appear.
Data & Asset Protection	Threats logged but nothing blocked in real time. No data lineage through AI pipelines. Sensitive data flows uncontrolled.	AI-specific DLP inline. Data lineage tracked. Inline prevention operational for prompts, data flows, and outputs.	Prevention at inference speed across all AI interaction points. Supply chain assessed before deployment.
Access & Execution Control	Multiple access strategies, none dominant. Enforcement varies by location. No runtime LLM guardrails.	Access policy follows users, devices, agents, and workloads regardless of location. Runtime guardrails embedded before go-live.	Access evaluated continuously against identity, context, and risk. Enforcement follows the workload regardless of location.
Detection & Response	WAFs degraded by AI traffic patterns. AI traffic indistinguishable from malicious activity to existing tools.	Deep payload inspection on AI traffic. WAF retuned for AI payloads. Model-layer filtering added.	AI-aware detection across all layers. Validated attack patterns automatically feed prevention rules.
Architectural Integration	Wide gap between strategy and architecture. Policy fragmented across environments. Tool sprawl adding complexity.	Consolidating onto platform architectures. Single policy layer spanning on-prem, cloud, and edge.	Single architecture under one policy engine and one console. AI-driven prevention across the full hybrid environment.

The progression across every capability points in the same direction: fewer disconnected controls, more shared visibility, and policy that can be enforced consistently across the hybrid environment. AI security maturity comes from turning governance, visibility, data protection, access control, and enforcement into one operating system, not from adding isolated protections one by one.

Five Actions to Close the AI Security Gap

AI security improves when organizations move from fragmented visibility and reactive controls toward unified policy, inline prevention, and architecture-level enforcement. These five actions are sequenced so each one creates the conditions for the next.

- 1 Build the AI asset inventory**
Map external AI services in use, internal AI applications and agents, model endpoints, credentials, and sensitive data flows. Visibility has to cover both workforce usage and enterprise AI systems.
- 2 Govern employee access to external AI services explicitly**
Apply consistent policy regardless of location, distinguish sanctioned from unsanctioned AI services, restrict personal or unapproved tenants where needed, and enforce data-handling controls on uploads and downloads.
- 3 Put prevention and runtime control into enterprise AI workflows**
For internal AI applications and agents, move controls into the prompt, data, output, and tool-execution path. Runtime validation and repeatable security testing need to be part of the application lifecycle, not an afterthought.
- 4 Give one function authority to define policy and prove enforcement**
AI governance becomes operational only when one accountable function can set policy, coordinate exceptions, and produce evidence that controls are actually being enforced across the environment.
- 5 Consolidate toward a unified hybrid security architecture**
Reduce fragmentation across datacenter, cloud, SaaS, endpoint, and AI-specific controls so visibility, detection, and prevention can work as one operating model rather than a series of disconnected point decisions.

Visibility makes everything else possible. Access governance reduces shadow AI and policy drift. Runtime and data-path controls reduce the highest-risk failure modes inside AI interactions. The full model becomes durable only when policy and enforcement can operate consistently across the hybrid environment.

Check Point Enables Hybrid Mesh Network Security

Together, every enforcement point can become an active sensor that feeds the AI-mesh intelligence, improving protection, cutting detection and prevention time from minutes to sub-milliseconds.



SECURING DATA CENTERS with Check Point Firewall

East-west traffic visibility
& enforcement

Zero-Trust workload-
to-workload access

High throughput with
ultra-low latency

AI Factory data centers



SECURING BRANCH AND APPLICATION CONNECTIVITY with Check Point SD-WAN

Advanced threat prevention

TLS inspection and
intrusion prevention

Policy-driven routing
decisions

Works alongside leading
SD-WAN vendors



SECURING THE CLOUD with Check Point Cloud Firewall

One consistent policy

Cloud-native visibility
& control

AI-powered
threat prevention

Auto-adjusting
dynamic policies



SECURING APPLICATIONS with Check Point WAF

Unified app security
across environments

Inline runtime prevention

API-first security



SECURING REMOTE USERS with Check Point SASE

Private & public access

Identity-first, continuous
access control

Device posture verification

Simplicity & visibility



UNIFIED MANAGEMENT with Check Point Services and AI Security

Identity-centric
management

End-to-end visibility
across the mesh

Cross-domain correlation
& risk prioritization

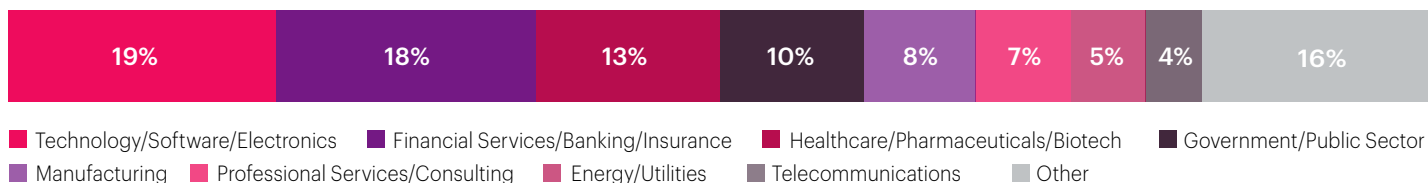
Architect for hybrid now and modernize in phases without losing control or consistency.

Methodology and Demographics

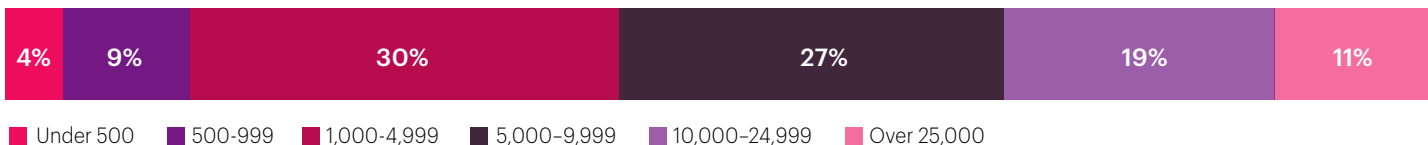
This AI security survey was conducted in early 2026 and gathered responses from 1,042 cybersecurity and IT professionals across a wide range of industries and organization sizes. Responses were collected through the Cybersecurity Insiders practitioner community of more than 600,000 security professionals worldwide. Respondents included CISOs, security architects, network engineers, security analysts, and IT leaders responsible for securing AI adoption, AI infrastructure, and the hybrid environments where AI now operates.

A stratified sampling approach ensured balanced representation across roles and segments, yielding a 95% confidence level with a $\pm 3.0\%$ margin of error to ensure valid industry representation.

INDUSTRY



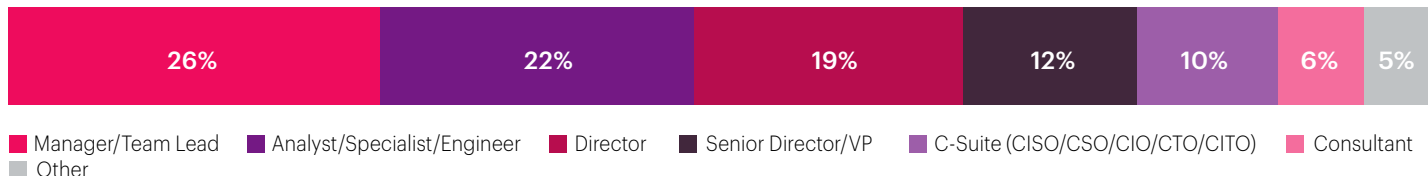
COMPANY SIZE



DEPARTMENT



CAREER LEVEL



©2026 Cybersecurity Insiders. All rights reserved.

Limited editorial citation (up to 100 words and one unaltered chart) is permitted with clear attribution to “Cybersecurity Insiders, 2026 Securing the AI Transformation Report” and a visible link to cybersecurity-insiders.com.

The report sponsor may reference the findings and use individual charts or data points in presentations and marketing materials with proper attribution. The full report, underlying dataset, and research methodology remain the intellectual property of Cybersecurity Insiders and may not be reproduced, redistributed, or incorporated into derivative research without written permission.

This report was produced by Cybersecurity Insiders with the support of **Check Point**. Permissions: info@cybersecurity-insiders.com

Cybersecurity

I N S I D E R S

BENCHMARK YOUR SECURITY MATURITY

Independent cybersecurity research revealing the gaps
that shape cybersecurity strategy

Cybersecurity Insiders produces independent research based on surveys of cybersecurity leaders and practitioners worldwide. Our reports reveal where security strategies break down in practice — helping organizations benchmark their maturity, identify capability gaps, and prioritize the actions needed to close them.

For more information, visit

[**cybersecurity-insiders.com**](https://cybersecurity-insiders.com)



Check Point Software Technologies Ltd. (www.checkpoint.com) is a global cyber security leader protecting more than 100,000 organizations worldwide. Its mission is to secure enterprises' AI transformation. With a prevention-first approach and an open ecosystem architecture, Check Point helps organizations block advanced threats, prioritize exposures, and automate security operations across complex digital environments. The unified architecture simplifies protection across hybrid networks, multi-cloud environments, digital workspaces, and AI systems. Structured around four strategic pillars, Hybrid Mesh Network Security, Workspace Security, Exposure Management, and AI Security, Check Point delivers consistent protection and visibility across multivendor environments, enabling organizations to reduce risk, improve efficiency, and accelerate innovation without increasing complexity.

www.checkpoint.com