

5 Blind Spots in *Your AI Footprint*

What every security leader should know about their organization's AI footprint

Introduction

In the last year, regular use of AI on corporate devices, authorized or not, has climbed from 15% to 45%¹. That is not a niche trend inside a few forward-leaning teams. That is roughly half of the workforce using AI as part of how they get work done, whether security signed off on it or not. And that number is continuing to rise.

The endpoint is where AI actually enters the organization. It is where employees adopt new tools, connect agents, and bring AI into the flow of their day, usually faster than any review process can keep pace with. And because agentic AI tools also inherit the permissions and access to corporate systems of the user, the endpoint is now also where the attack surface is growing the fastest.

First, AI introduces new types of security risks regarding what AI is allowed to do on the endpoint and how it connects to other systems. But it has also made existing endpoint blind spots much more urgent. For years, security teams have lived with challenges around incomplete inventory, stale policies, and unmanaged configurations in software applications. Malicious actors now have access to frontier, Mythos-class models that can find and exploit these gaps at machine speed. Security leaders need to act fast to reduce their risk exposure.

¹ Verizon Business, “2026 Data Breach Investigations Report”

The five questions in this eBook are intended to help surface common challenges related to the safe adoption of AI in the enterprise.

They should work as a diagnostic, not a checklist to pass or fail.

Any gaps are worth finding now, on your own terms, rather than during an incident review.

The 5 blind spots

01

Do you know where AI is actually being used across your organization, not just where you've approved it?

02

Of the tools you've approved, does real-world usage still match your policy?

03

Are your approved AI tools still operating within your security guardrails?

04

Do you have visibility into what your AI tools connect to and depend on?

05

If you find something risky, can you act on it without slowing the business down?

01

Do you know where AI is actually being used across your organization, not just where you've approved it?

Most employees are not trying to evade policy. They are trying to work faster, and AI tools offer immediate value. The problem is that unmanaged adoption is more consequential than traditional shadow IT. An agentic tool may be able to access files, credentials, development environments, and business systems through the permissions of the person using it. The more autonomy it has, the more it can do before security knows it is present.

The cost of not knowing

Personal and free-tier accounts may provide weaker administrative and data protections. Sensitive company information can move outside approved controls or be handled under terms the organization has not reviewed.

Security cannot reliably confirm that employees are using an approved account, a vetted version, or a fully patched release.

A compromised or misconfigured agentic tool can act with real user permissions, increasing the potential impact beyond simple data exposure.

What good looks like

Continuous discovery of the AI tools, accounts, extensions, local models, and connections actually present on endpoints. Point-in-time surveys and login data alone are not enough because they miss local tools, unauthenticated use, and changes between reviews.

02

Of the tools you've approved,
does real-world usage still match
your policy?

AI policy is being written while the category is still changing. New tools appear quickly, approved products gain new capabilities, and usage expands into new teams and data types. A tool reviewed six months ago may now be used by more people, for different workflows, and with broader access than the original approval anticipated.

As a result, policy can become stale even when the underlying intent remains sound. The challenge is not simply writing a policy. It is keeping policy aligned with how AI is actually being used.

The cost of getting this wrong

Policies built without a current risk picture tend to be too broad, creating unnecessary restrictions, or too narrow, missing the tools and behaviors that carry the greatest exposure.

Adoption can outpace review. A policy based on last quarter's tools may not address the applications and capabilities employees are using today.

Blanket blocking can disrupt legitimate work and encourage employees to move to personal accounts or other unsanctioned alternatives.

What good looks like

A living inventory that connects organizational policy to observed behavior. Security teams should be able to see where usage differs from policy, understand the associated risk, and refine guardrails without lengthy rule-building or change processes. This supports a risk-based approach that protects the business without defaulting to a total block.

03

Are your approved AI tools still operating within your security guardrails?

Approval is a point-in-time decision. AI products continue to change after procurement, often through new features, updated defaults, and expanded integrations. A capability may be enabled automatically, a data setting may change, or a new connection may become available without a formal security review.

Control also depends on where settings are enforced. Central administration for many AI tools is still developing, and some controls can be changed locally by users. An approved product can therefore drift outside the configuration security originally evaluated.

The cost of assuming approved means safe

Vendor updates can change the risk profile of a tool even when no one inside the organization changes the configuration.

Local overrides may re-enable data sharing, broaden permissions, or introduce new integrations without creating a visible exception for security.

Configuration reviews often happen at procurement and are not repeated until renewal or an incident, allowing drift to accumulate unnoticed.

What good looks like

Continuous monitoring of security-relevant configurations across approved AI tools. Changes to permission scopes, data handling, execution privileges, and connected services should be identified when they move outside policy, not months later.

04

Do you have visibility into what your AI tools connect to and depend on?

The AI application is only the visible front end. Extensions, plugins, packages, models, and connected services form a broader ecosystem behind it. These components can expand what the tool can access and do, yet many existing security programs were not designed to inventory or assess them.

Dependencies also change independently. A plugin can update, request additional permissions, or be replaced without a major change to the parent application. A clean review of the application does not guarantee that the surrounding ecosystem will remain safe.

The cost of stopping at the app

The primary application may appear trustworthy while one of its dependencies introduces the actual risk.

Component marketplaces often have inconsistent verification and testing, making impersonation, abandoned packages, and malicious components harder to distinguish from legitimate ones.

Dependencies can update or change after approval, reducing the value of a one-time assessment.

Traditional vendor risk and software supply chain processes may not cover the local components and connections employees add themselves.

What good looks like

Continuous mapping of the components, permissions, and connections behind each AI tool. Security teams need to understand not only which application an employee opens, but also what that application depends on and what those dependencies can access.

05

If you find something risky, can you act on it without slowing the business down?

Visibility is necessary, but it does not reduce exposure on its own. If every finding becomes an alert, ticket, and manual investigation, security remains reactive and the queue grows faster than the team can resolve it.

Effective control requires more than detection. Security teams need targeted ways to restrict, reconfigure, remove, or block software based on organizational policy. The level of automation should match the risk, from human approval for sensitive actions to autonomous remediation for well-defined issues.

The cost of visibility without action

Known risks remain open while they wait for manual triage and coordination across security, IT, and end users.

As AI adoption accelerates, the gap between identifying an issue and fixing it can continue to widen.

When targeted remediation is difficult, teams are more likely to use broad blocking, creating avoidable business disruption.

Without a preventive control behind the fix, the same issue can return when software is reinstalled, reconfigured, or adopted by another employee.

What good looks like

Policy-based controls that can allow, restrict, reconfigure, remove, or block software with the appropriate level of human oversight. Remediation should keep pace with the rate of change, and preventive policies should stop resolved issues from repeatedly returning.

Where does your org land?

Few security leaders can answer all five questions with complete confidence. That is not a failure of the security team. It reflects how quickly AI has entered the enterprise and how little of the existing security stack was designed to govern this new layer of endpoint activity.

Progress starts with a trusted source of truth for what is actually running across the environment. That foundation makes practical governance possible because policy can be tied to observable behavior. With governance in place, security teams can move from discovering risk after the fact to preventing known issues from becoming exposure.

This progression, from trusted visibility to governance to prevention, is the approach Glow uses to help security teams take control of their AI footprint.

See how your organization measures up across all five blind spots.

Learn more and request a demo at glow.io

