

FOS

State of the Internet | V12 Issue 04 | 2026



Enterprise AI Usage Risk Report 2026



Contents

03	Introduction
03	A practical, fact-based analysis of the top 5 enterprise AI risks of 2026
04	The expansion of the threat surface via the proliferation of AI in the workplace
06	Security spotlight: The rise of AI power users is creating a new enterprise risk surface
07	Personal AI accounts and unapproved AI apps that drive shadow AI
10	Security spotlight: When an unapproved AI tool becomes a national security concern
12	Data leak risks via AI data exposure
13	Security spotlight: Executive workflow abuse and public LLM data spills
14	Browser and IDE extensions that go unnoticed in most organizations
17	Security spotlight: The anatomy of CursorJacking
18	AI agents that penetrate the enterprise environment and operate outside of existing guardrails
19	Security spotlight: CometJacking — When attackers target the AI instead of the user
20	Mitigation framework: How CISOs can reduce enterprise AI risk
24	A practical checklist for CISOs and security leaders
26	Conclusion

Introduction

When the State of the Internet/Security report first analyzed the emerging [artificial intelligence \(AI\)](#) landscape in early 2025, enterprise adoption was largely characterized by caution and curiosity; today, it is a structural mandate. What began as sporadic experimentations with ChatGPT has evolved into a sprawling ecosystem of AI assistants, AI browsers, [AI agents](#), AI extensions, and autonomous workflows that touch every aspect of enterprise work.

Employees use AI to write code, analyze contracts, summarize meetings, generate content, research customers, and automate repetitive tasks. As a consequence, enterprises are racing to capture the productivity benefits of AI while simultaneously struggling to understand where their data goes, who has access to it, and how these systems can be abused.

Unlike previous technologies, AI continuously consumes, generates, stores, and acts upon enterprise data. The downside is that AI introduces entirely new categories of risk that traditional security tools were never designed to address.

With Akamai's recent acquisition of LayerX, we received access to their data. As a result, this special edition State of the Internet (SOTI) examines how organizations are adopting AI, where governance is failing, and how attackers are already exploiting AI-powered workflows.

For CISOs and security leaders, this report dissects where (and how) AI is being used in the enterprise and how it can be exploited by threat actors. It also provides a roadmap to limiting the risk from these attack vectors without losing the productivity enhancements of AI.

A practical, fact-based analysis of the top 5 enterprise AI risks of 2026

This report focuses on the top five risks from AI use that enterprises face today:

1. The expansion of the threat surface via the proliferation of AI in the workplace
2. Personal AI accounts and unapproved AI apps that drive shadow AI
3. Data leak risks via AI data exposure
4. Browser and IDE extensions that go unnoticed in most organizations
5. AI agents that penetrate the enterprise environment and operate outside of existing guardrails

The expansion of the threat surface via the proliferation of AI in the workplace

AI has become a mainstream enterprise technology, with employees increasingly using AI assistants to research information, generate content, write code, analyze data, and automate routine tasks.

Although AI adoption is widespread, usage is highly concentrated among a relatively small group of “AI power users” who rely on AI throughout their daily workflows. These AI power users take part in significantly more conversations, engage in longer interactions, and increasingly use AI as a collaborator rather than as a simple search tool. These users are also more likely to share business information, upload files, submit sensitive data, and integrate AI into day-to-day decision-making.

AI adoption vs. usage concentration

According to their [State of AI Usage Report 2026](#), LayerX observed that across their platform, nearly 20% of employees used AI weekly, while more than 30% used it monthly. Slightly more than 40% engage with AI quarterly, demonstrating that AI has expanded well beyond technical teams into everyday business functions. However, the report showed that most users are not daily or weekly users; almost 50% use AI more casually (Figure 1).

Enterprise Usage of AI Tools Across Periods

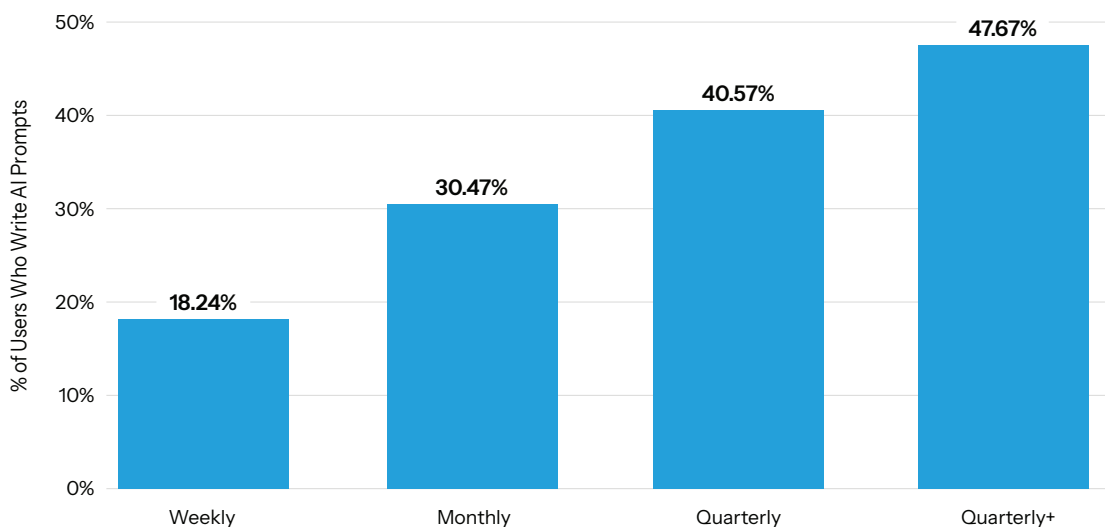


Fig. 1: Nearly 20% of employees use AI weekly

The LayerX data also indicated that the average enterprise user participates in 36 AI conversations, yet the bottom 50% of users engage in 12 conversations or fewer. In contrast, the top 5% of users generate at least 144 conversations, highlighting the emergence of a distinct group of AI power users (Figure 2).

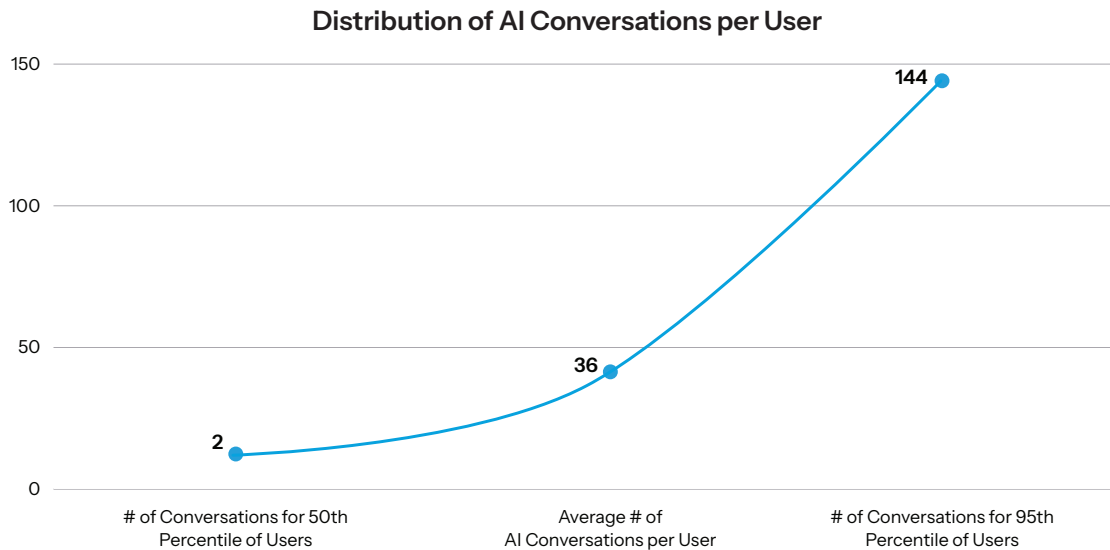


Fig. 2: While the average employee participates in 36 AI conversations, the bottom 50% of casual users engage in 12 or fewer AI conversations

The same pattern appears in conversation depth. The average AI conversation contains approximately five prompts, suggesting that many interactions are simple, one-time requests. However, the top 5% of conversations contain at least 18 prompts, indicating that some employees are increasingly using AI as an iterative collaborator rather than as a simple question-and-answer tool (Figure 3).

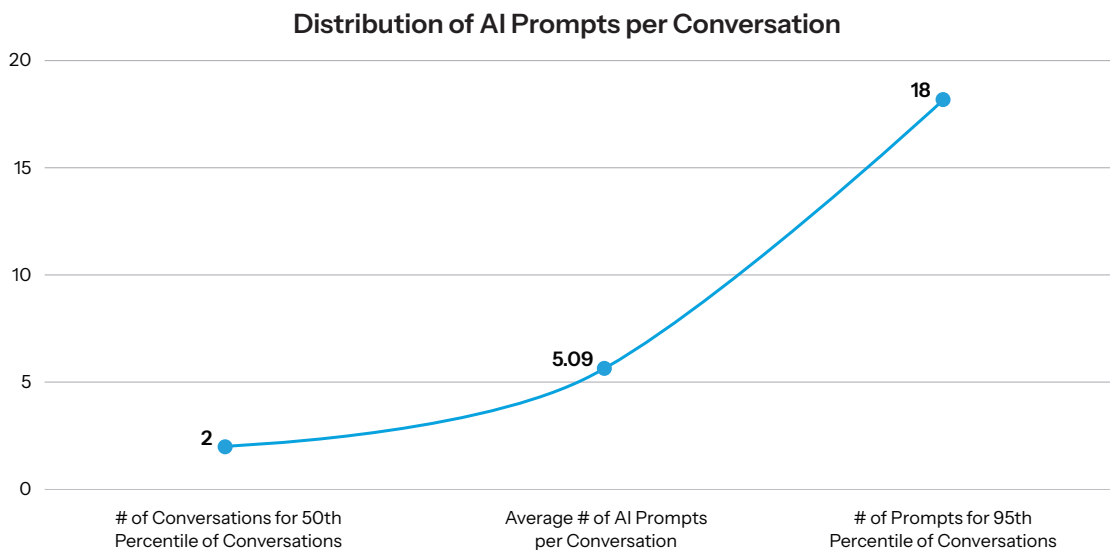


Fig. 3: The top 5% of AI conversations exceed 18 prompts, reflecting a trend toward iterative collaboration over simple questions and answers

An interesting pattern emerged when LayerX compared employee activity and conversation depth: Both graphs follow almost identical distributions. In each case, a small percentage of users account for a disproportionate amount of AI activity. This reinforces that enterprise AI risk is driven by a concentrated group of heavy users rather than by the workforce as a whole.

Security spotlight

The rise of AI power users is creating a new enterprise risk surface

Just like pen and paper, a laptop, and a mobile phone, AI tools are now part of the standard enterprise kit. This includes AI assistants such as [ChatGPT](#), [Microsoft Copilot](#), [Google Gemini](#), and other [frontier LLMs](#), as well as highly integrated autonomous coding tools.

While these productivity tools are quickly becoming essential, they also expand an organization's threat surface for exploitation.

In 2026, LayerX researchers demonstrated how a popular frontier AI coding assistant could be manipulated through an adversarial prompting technique dubbed [vibe hacking](#).

The demonstration showed that by subtly modifying a project's local markdown instruction file (such as an [AL_CONFIG.md](#) or similar behavioral profile file used to guide the model's contextual boundaries), an attacker could covertly influence how the AI assistant operated inside a development environment.

The manipulated instructions could cause the frontier AI assistant to perform unauthorized actions or generate insecure outputs that aligned with the attacker's objectives — all while appearing to be part of the developer's normal automated workflow.

The significance of this tactical risk extends far beyond software development. Organizations are increasingly seeing the emergence of AI power users who rely on [large language models \(LLMs\)](#) throughout their daily workflows and generate significantly more AI interactions than the average employee.

These AI power users often:

- Conduct longer, deeply contextual AI sessions
- Share highly sensitive business context with AI systems
- Delegate execution-level work to autonomous AI agents
- Depend heavily on AI-generated outputs to make operational decisions

As AI becomes embedded into critical workflows, manipulating the underlying frontier model can become just as valuable to attackers as compromising the user directly. The implications of vibe hacking means organizations no longer need to ask *if* employees use AI. Instead, they must identify who depends on it, how deeply frontier models are embedded in daily workflows, and where the highest concentration of risk may lie.

Personal AI accounts and unapproved AI apps drive shadow AI

While major platforms like ChatGPT, Microsoft Copilot, and Google Gemini dominate the headlines, they represent only the visible surface of the enterprise AI landscape. For many organizations, AI governance remains heavily focused on managing these few well-known platforms — creating a significant visibility gap with regard to the true scale of decentralized AI adoption across the company.

The LayerX data shows that while the top four or five AI applications in most organizations are widely used, there is a remarkable drop-off in usage after that. The top four enterprise AI applications are used, on average, by more than 20% of the organization, but after the 10th AI application that figure drops to less than 5% and reaches nearly 0% by the 30th application (Figure 4).

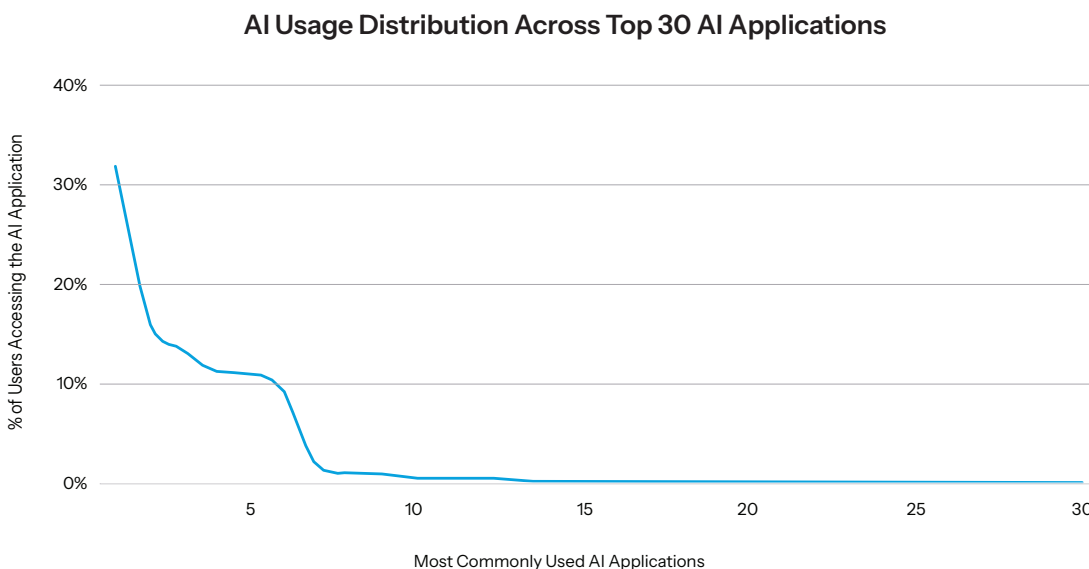


Fig. 4: The percentage of use of the top 30 AI applications in enterprise environments

The implications of this data pose a series of interesting questions for organizations. While most security professionals can easily name the top four or five AI applications running on their network, a vast “long tail” of [shadow AI](#) apps operates just below the surface — often used by small, isolated teams without corporate oversight. These can be niche applications for specific purposes, region- or language-specific applications, or personal preferences of a handful of people, but to which the enterprise has no visibility, governance, or oversight.

Another significant enterprise security challenge identified by LayerX is the growing disconnect among corporate identities, personal AI subscriptions, and unmanaged AI services. As with mobile devices, employees increasingly “bring their own AI tools (BYOAI)” and access AI through personal accounts that operate outside enterprise security controls, creating additional visibility gaps around how business data is stored, retained, and processed.

The result is shadow AI; that is, business activity occurring in AI environments that IT never approved, security teams can't monitor, and compliance teams can't audit. [Nearly half of all enterprise AI conversations](#) (47.11%) take place through personal identities rather than corporate-managed accounts (52.89%; Figure 5).

Identity Distribution in AI Applications: Corporate vs. Personal

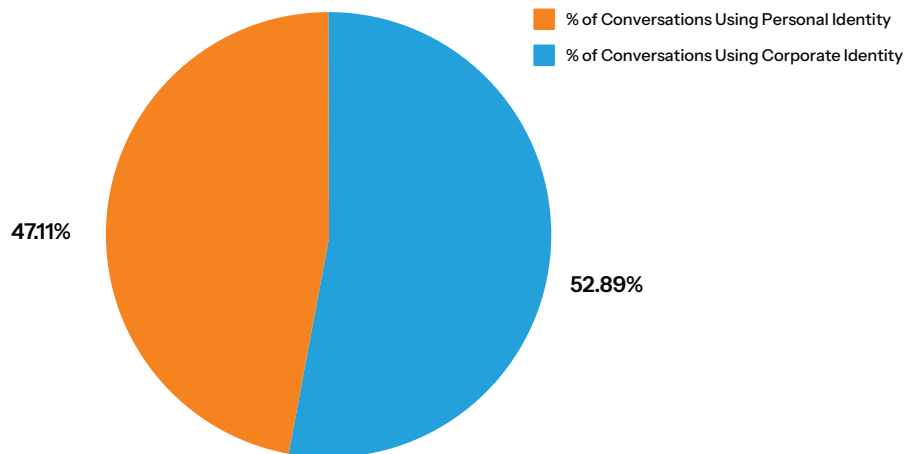


Fig. 5: Nearly half of all enterprise AI conversations occur via unmanaged personal identities, creating a major governance visibility gap

Personal AI subscriptions create a false sense of governance: Although activity appears to be corporate, it often falls outside enterprise security, compliance, and audit controls. Organizations may be unable to verify data residency, enforce retention and deletion policies, demonstrate regulatory compliance, or ensure that sensitive business information is handled according to corporate and contractual requirements.

One of the most surprising findings LayerX observed was that 14.4% of enterprise AI conversations occurred using corporate email addresses linked to personal “freemium” AI subscriptions rather than to enterprise-managed licenses (Figure 6). This means that even though they were being used via a corporate email address, the data going into those accounts was being used by the LLM for public model training.

Conversations Using Corporate Identities: Enterprise vs. Personal License

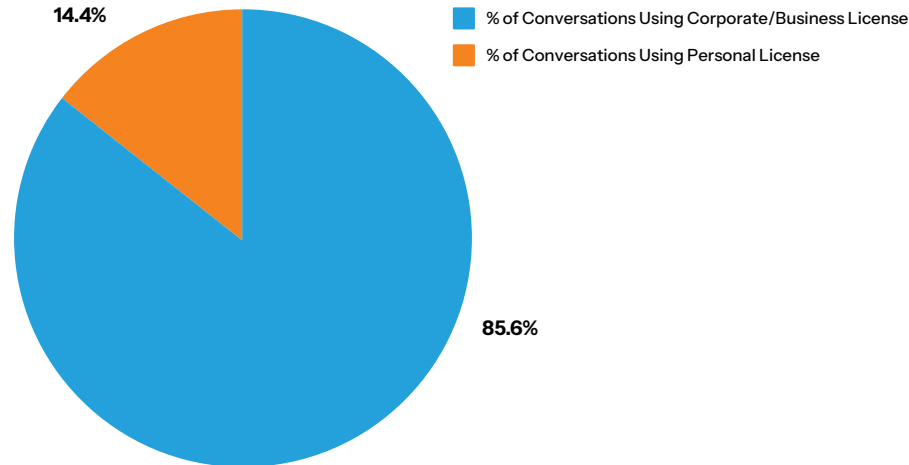


Fig. 6: A total of 14.4% of corporate email AI conversations are linked to personal subscriptions that bypass enterprise compliance controls

Moreover, the LayerX platform observed that the distribution of personal and corporate use varies significantly by platform. ChatGPT, Copilot, and DeepSeek are primarily accessed through personal accounts, with more than 60% of conversations tied to personal identities. In contrast, enterprise-focused offerings such as Gemini Enterprise and Copilot M365 are overwhelmingly used through corporate-managed accounts, representing 98.15% and 90.55% of conversations, respectively (Figure 7).

Corporate vs. Noncorporate Identity by Platform

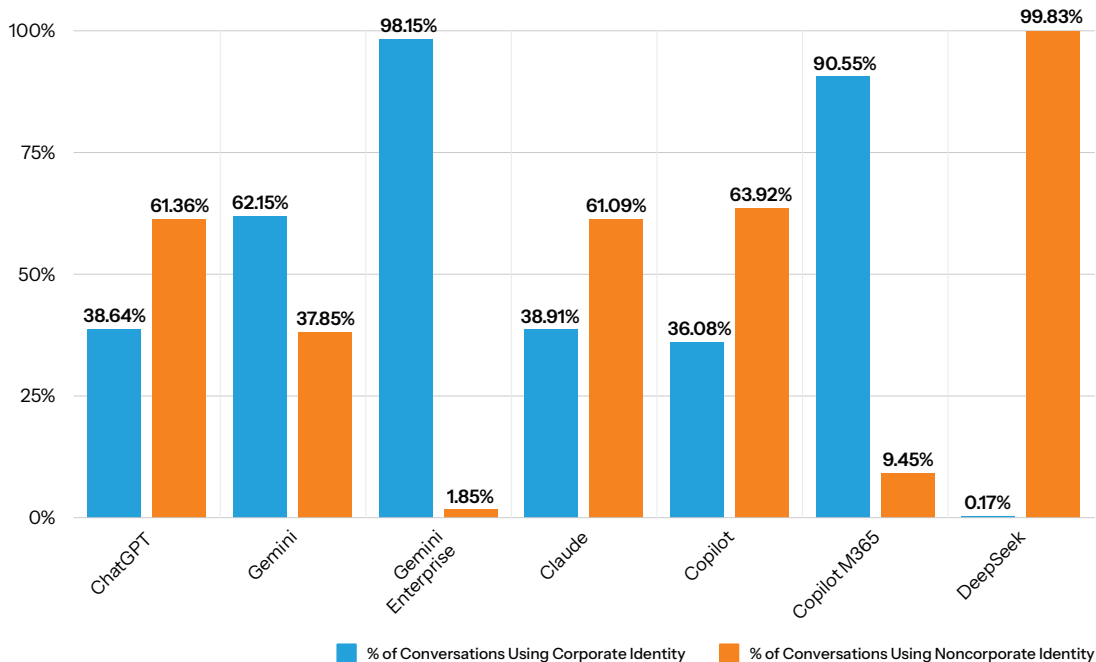


Fig. 7: Corporate identity use varies significantly across platforms

When an unapproved AI tool becomes a national security concern

In 2026, [U.S. policymakers intensified scrutiny of Chinese AI company DeepSeek](#), with proposals ranging from government restrictions to adding the company to the [U.S. Entity List](#). Concerns centered on data security, potential foreign influence, and the risks associated with sensitive information being processed by AI systems outside U.S. jurisdiction.

The debate emerged just months after DeepSeek experienced explosive global adoption and became one of the fastest-growing AI applications in the world. This case perfectly illustrates the friction between instant, user-driven software adoption and critical data security boundaries.

DeepSeek's explosive adoption pattern follows a familiar, risky trajectory for enterprise security (Figure 8).

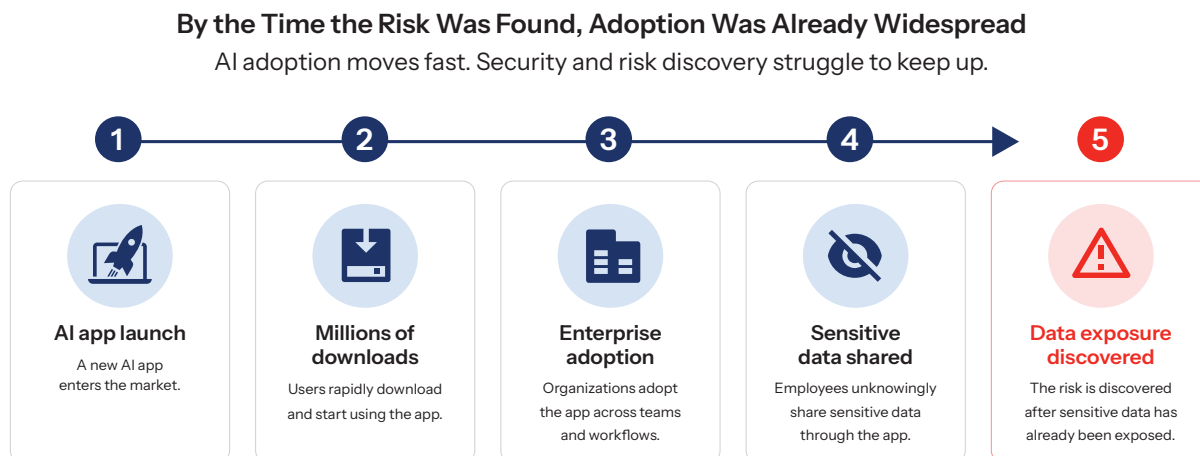


Fig. 8: The gap between AI adoption and risk detection

First, employees discover and adopt a new AI capability, which quickly scales across teams and begins consuming sensitive business data. Consequently, corporate security and governance reviews are forced into a reactive posture — occurring only after the platform has already been integrated into the daily workflow.

However, the operational challenge extends far beyond any single platform. As increasing numbers of employees use multiple niche tools — frequently bypassing corporate credentialing through the use of personal accounts — organizations suffer a severe degradation in visibility. Security teams are increasingly unable to verify:

- **Inventory:** Which specific AI tools are actively running on the corporate network
- **Data lineage:** What classified or proprietary business information is being transmitted
- **Data residency:** Where that ingested data is physically stored or cloud stored
- **Lifecycle management:** How long the vendor retains the data and whether it is being absorbed into public model training pipelines

The compounding effect of these visibility gaps is a rapidly expanding, unmanaged AI ecosystem that introduces systemic security, privacy, and regulatory compliance risks across the enterprise.

The friction points surrounding emerging tools demonstrate that AI governance can no longer focus solely on a handful of approved enterprise assistants. The strategic imperative for modern CISOs is to establish continuous visibility and control over the growing, decentralized long tail of AI applications that employees can adopt independently and deploy instantly.

Data leak risks via AI data exposure

Traditional data loss prevention (DLP) was designed for a different era of enterprise computing: It assumes that sensitive information moves through relatively well-defined channels, such as email, file transfers, cloud storage, or web uploads. Detection typically relies on exact matching, fingerprints, regular expressions, dictionaries, and predefined data classifications.

AI fundamentally changes that model.

Instead of moving complete files, employees increasingly share information as unstructured data, through fluid channels such as:

- Prompts
- Conversational context
- Snippets of source code
- Screenshots
- Copied text
- Iterative conversations
- Generated responses

As a result, sensitive information has become fragmented across dozens of interactions that individually may appear harmless but collectively expose valuable business context. AI creates an entirely new class of data movement and risk that legacy controls were never designed to understand.

The biggest AI security risk is no longer employees who access AI, it is employees who share sensitive business data with AI.

The findings point to a broader reality: AI data leakage is fundamentally different from traditional data leakage. Protecting enterprise data now requires understanding the context of every AI interaction and not simply whether data matches predefined patterns.

Executive workflow abuse and public LLM data spills

In early 2026, [multiple news outlets](#) reported that a U.S. government official accidentally released internal, restricted operational data from the Cybersecurity and Infrastructure Security Agency (CISA) via a public AI infrastructure. The event forced a macro-level review of how federal agencies interact with commercial third-party models and how sensitive information was being shared with external AI systems. Crucially, the incident was not an external compromise; it was an optimization mistake made by a highly trained user trying to execute routine daily analysis.

The exposed information included:

- Government documents intended for internal use
- Sensitive operational information
- Data that was not classified, but was not intended for public AI systems

As noted, the incident was not the result of a cyberattack or technical compromise. Instead, it reflected a growing pattern seen across organizations: Employees increasingly use AI tools to help complete everyday tasks and may upload documents, paste information, or share data without fully considering where that information is being processed or stored (Figure 9).

Even Security Experts Can Leak Data

AI data leaks are often the result of routine work rather than malicious intent

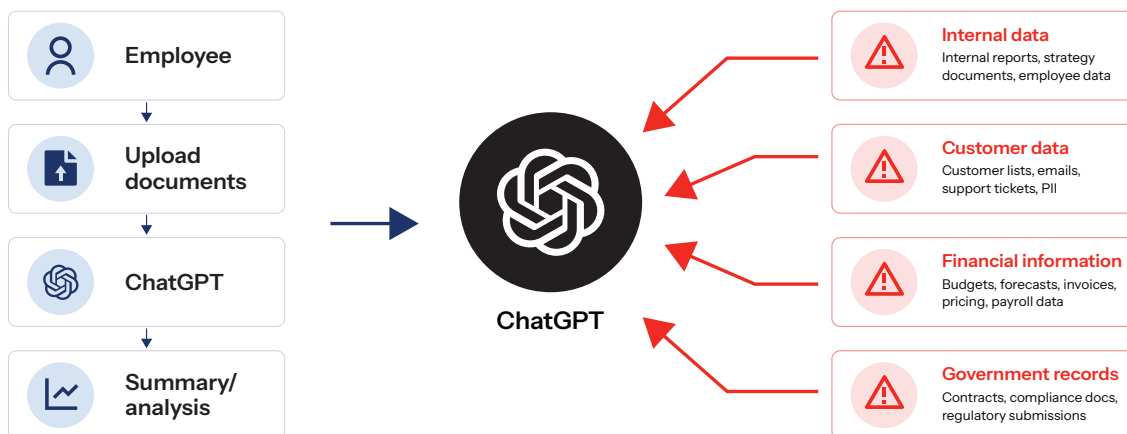


Fig. 9: Routine AI use can lead to unintentional data exposure

Enterprise data security has fundamentally shifted. Sensitive data no longer leaks solely through traditional vectors like email attachments or file transfers; it flows directly into AI inputs. Security teams must classify AI interactions as a brand-new category of data movement, applying the exact same visibility, governance, and protection controls used for any other critical corporate asset.

Successfully securing AI is no longer about deciding whether employees can use AI. Instead, it requires organizations to establish a new operating model built around five principles:

1. Employees understand what data can and cannot be shared.
2. Every AI interaction is visible.
3. Sensitive information is evaluated in context, not just by pattern matching.
4. Controls adapt to user risk, application risk, and data sensitivity.
5. AI remains productive while risky interactions are prevented in real time.

Browser and IDE extensions that go unnoticed in most organizations

Most organizations focus their AI governance efforts on web-based AI applications such as ChatGPT, Copilot, and Gemini. However, employees are increasingly accessing AI through **browser extensions** that operate directly inside the browser and often have extensive access to web content, user activity, credentials, and enterprise applications.

Unlike stand-alone AI applications, browser extensions run continuously alongside users and can access sensitive data across multiple websites and workflows. They create a largely unmanaged AI access channel.

AI extension use is emerging across enterprises, with small and medium-sized enterprises showing the highest adoption rate. Almost 18% of users at medium-sized enterprises run at least one AI extension, according to [LayerX](#). Their research also identified that larger organizations remain more conservative or regulated, with 9.53% of users running an AI extension (Figure 10).

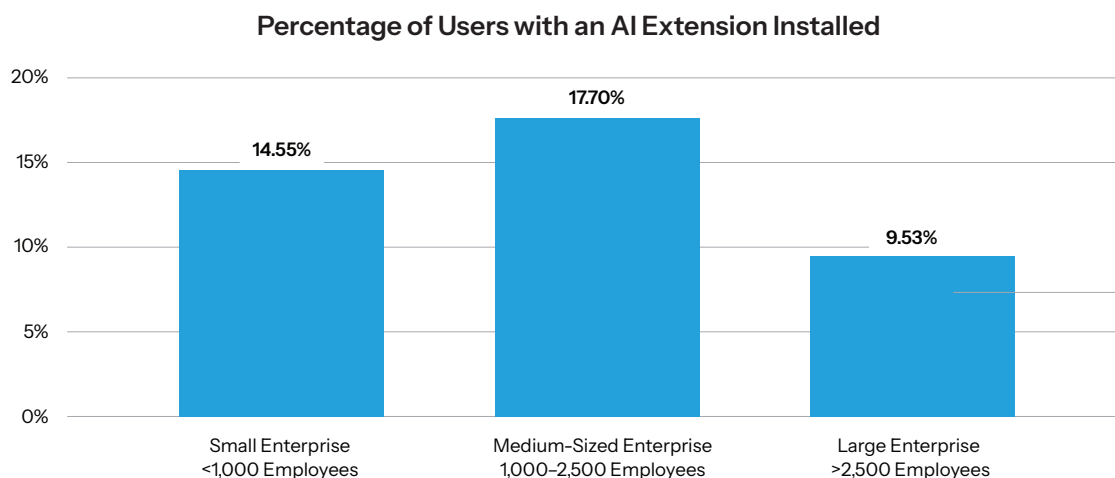


Fig. 10: Close to 18% of users at medium-sized enterprises run at least one AI extension, compared with 9.53% in larger, more heavily regulated organizations

The [LayerX research](#) also indicated that nearly 75% of AI browser extensions request high or critical permission levels (56.4% high, 16.7% critical), while only 1.4% operate with low permissions (Figure 11).

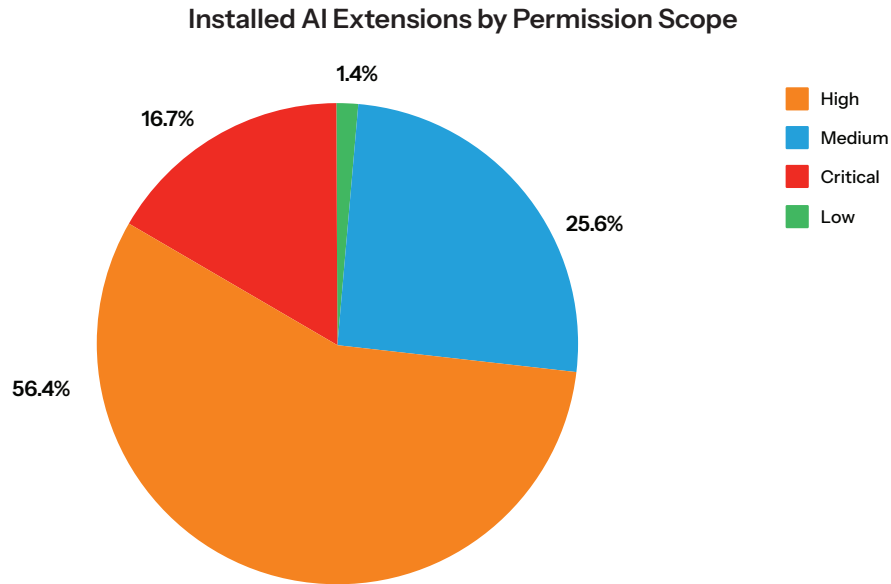


Fig. 11: Nearly 75% of AI extensions demand high (56.4%) or critical (16.7%) permissions, exposing sensitive browser data and user sessions to potential exfiltration or hijacking

Extensions with elevated permissions can access sensitive browser data and user activity, which means a malicious or compromised extension could easily expose sensitive information or take over user sessions.

AI extensions show a higher security risk profile: 16.31% of AI extensions have known CVEs, compared with 10.8% across all extensions (Figure 12). The higher CVE rate in AI extensions means that they require stricter vetting, monitoring, and management before being deployed in enterprise environments.

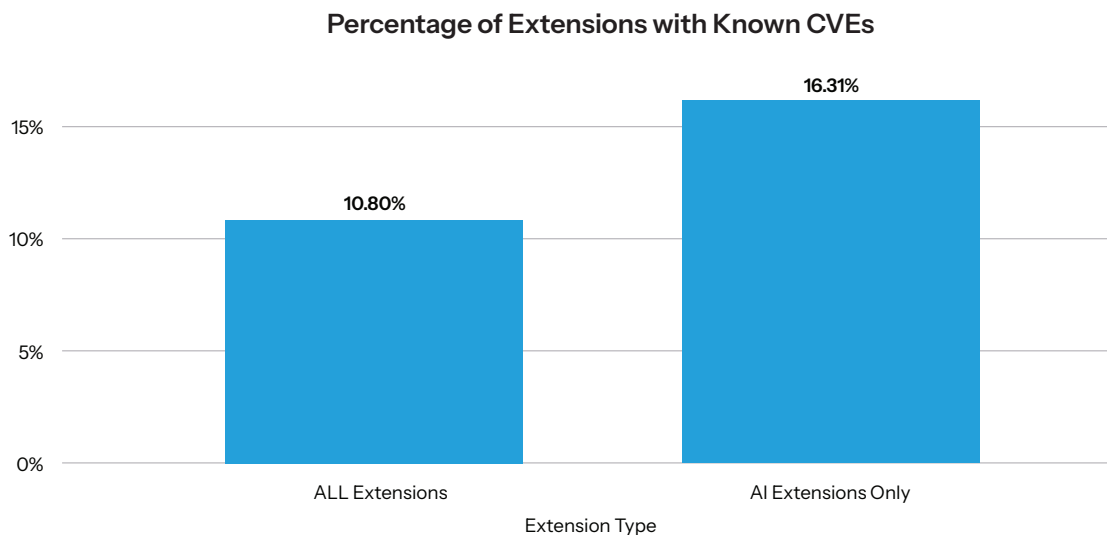


Fig. 12: More than 16% of AI extensions contain known CVEs compared with 10.8% of standard extensions

Moreover, AI extensions are nearly three times more likely to request cookie access than the average browser extension (Figure 13). Also, 41.91% of AI extensions request scripting access, compared with 15.4% of all extensions, according to the LayerX report findings (Figure 14).

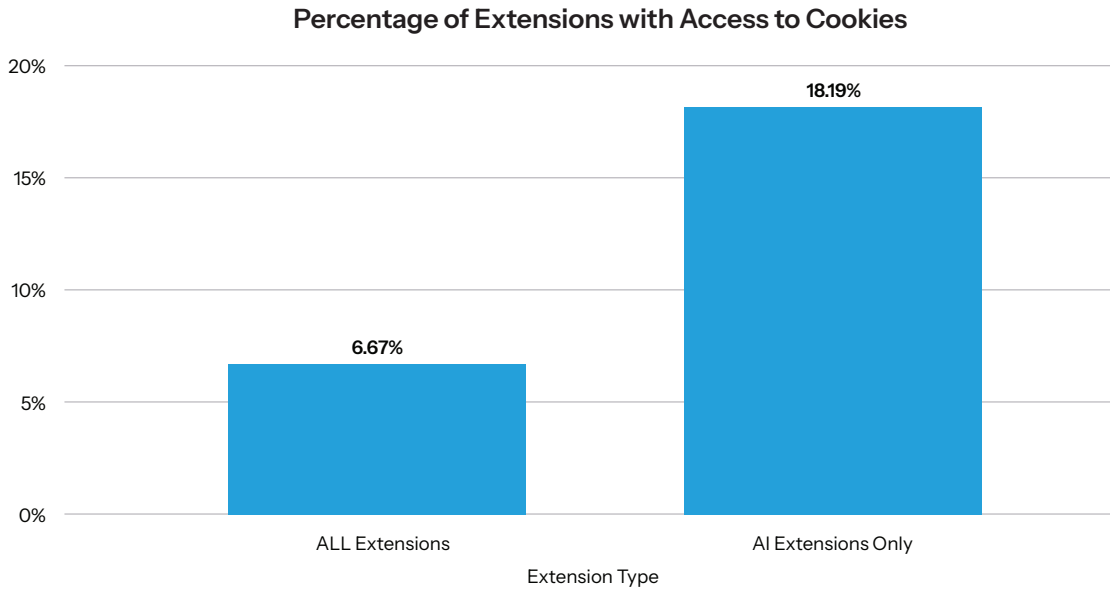


Fig. 13: AI extensions are almost 3x more likely to request cookie access vs. average browser extensions

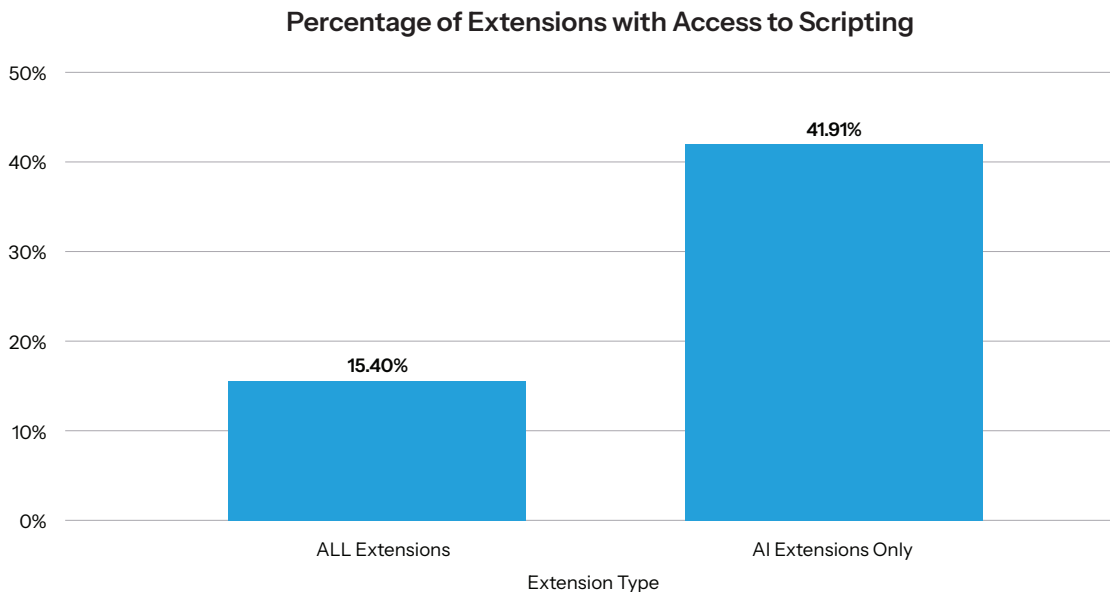


Fig. 14: Nearly 42% of AI extensions request powerful scripting access — compared with just 15.4% of standard extensions — which introduces severe privacy and security risks

AI extensions introduce disproportionately higher privacy, security, and data-access risks and require stricter governance, review, and ongoing monitoring in enterprise environments.

The anatomy of CursorJacking

In 2026, researchers at LayerX [uncovered a critical security flaw that they dubbed “CursorJacking”](#) — a novel attack vector against users of the popular Cursor AI coding assistant.

Their findings demonstrated that a rogue browser extension could silently exfiltrate Cursor API keys directly from the browser environment. By weaponizing the broad permissions often granted to everyday extensions, an attacker could seamlessly impersonate developers and compromise their entire AI-driven development environment (Figure 15).

Attack Flow: Malicious Cursor Extension Steals Credentials

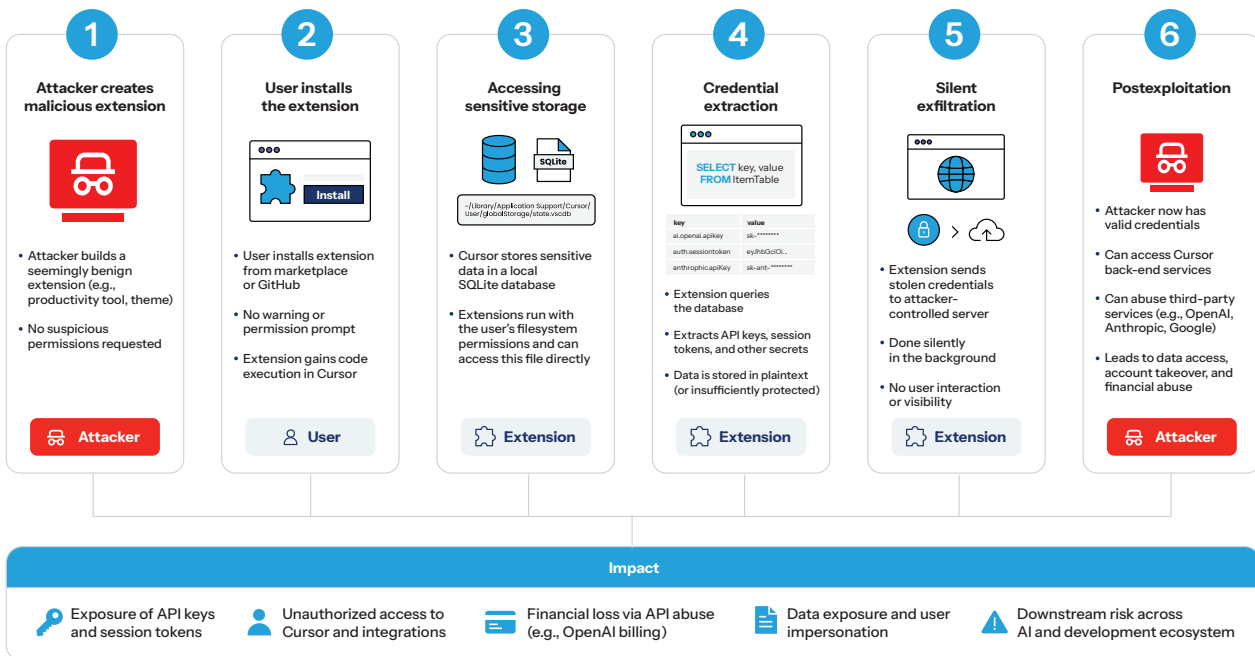


Fig. 15: The anatomy of an AI extension credential theft attack

If exploited, a CursorJacking attack could expose highly sensitive assets, including:

- Cursor API keys: Granting unauthorized access and potential cost inflation
- Proprietary source code: Intercepting the active code context shared with the assistant
- AI conversations and prompts: Revealing proprietary logic, intellectual property, or internal queries
- Downstream workflows: Compromising connected repositories and developer pipelines

The CursorJacking research underscores a critical evolution in enterprise identity risk. Because AI-powered extensions operate directly within the browser or development environment, they possess broad, programmatic access to a user's entire AI workflow. This includes real-time user prompts, proprietary code bases, AI-generated outputs, and downstream connected cloud services. If an attacker successfully compromises or subverts these extensions to harvest active AI session credentials or API keys, they gain a persistent foothold into both the application itself and the broader corporate workflows connected to it.

CursorJacking demonstrates that AI extensions have become high-value conduits for enterprise exploitation. To mitigate this risk, security teams must treat extensions as the primary gateway between users and sensitive corporate data. This requires absolute visibility into extension inventory, active permissions profiles, credential storage mechanisms, and overall data access boundaries.

AI agents that penetrate the enterprise environment and operate outside of existing guardrails

AI is no longer limited to answering simple questions. Enterprises are rapidly deploying autonomous AI agents that can make independent decisions, take direct actions, and execute complex, multistep workflows with minimal human oversight. Simultaneously, AI browsers and agentic interfaces are fast becoming the primary gateways through which users interact with corporate systems.

As these agents gain deep access to enterprise [software as a service \(SaaS\)](#) applications, critical business data, and core operational workflows, they transition from passive productivity tools into active participants within the organization — leaving security teams desperate for granular visibility into both human and agent interactions.

Security spotlight

CometJacking — When attackers target the AI instead of the user

In 2025, [LayerX researchers](#) analyzed a new attack technique called “CometJacking” against Perplexity’s Comet AI browser.

By embedding malicious instructions within a web page, researchers showed that an attacker could manipulate the browser’s AI agent through an indirect prompt injection attack.

The AI agent could potentially:

- Expose sensitive information
- Access browser content and emails
- Interact with credentials and local files
- Perform actions the user never intended

All from simply visiting a web page (Figure 16).

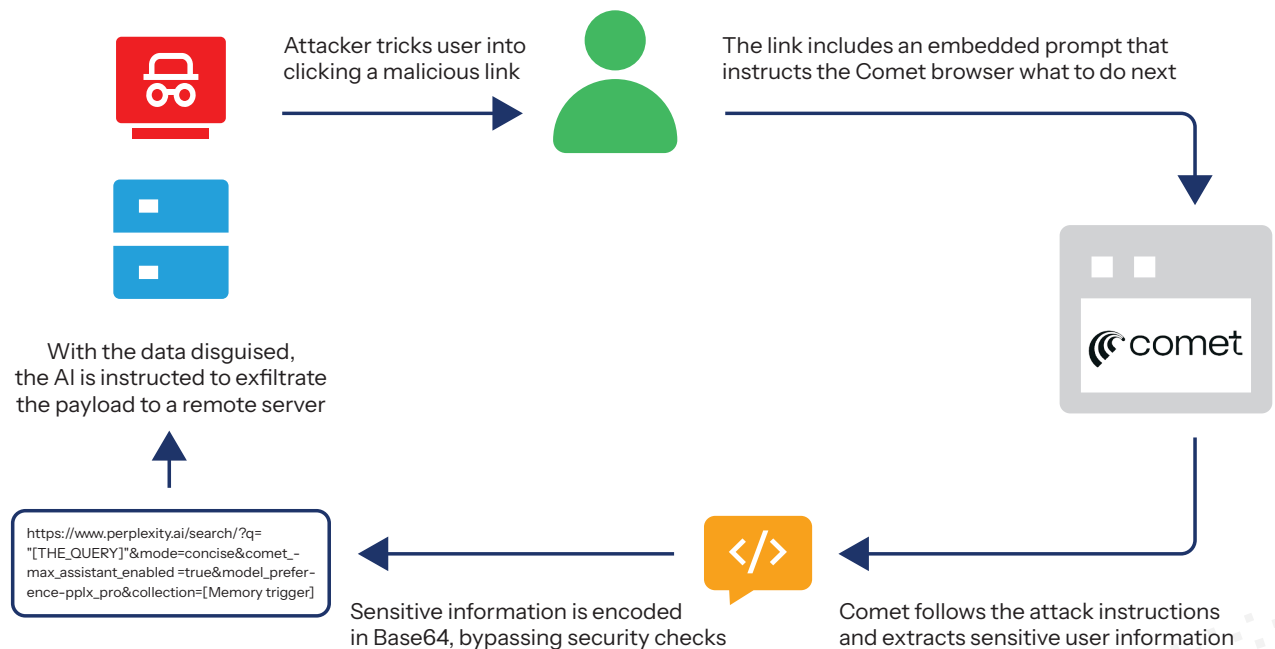


Fig. 16: The anatomy of a prompt injection attack against an AI browser

The CometJacking exploit highlights a fundamental shift in enterprise identity risk. Historically, cybercriminals targeted the human user via phishing links or compromised web applications. With the rise of AI-native browsers, adversaries are pivoting to target the digital agent that is acting on the user's behalf. Because these browsers are granted deep access to enterprise SaaS environments, email portals, calendars, and active session credentials, they effectively function as powerful, highly privileged intermediaries.

A successful exploit against the browser's AI engine allows an attacker to bypass traditional user authentication entirely — manipulating the agent to exfiltrate sensitive local data or connected files directly.

The CometJacking attack technique demonstrates that organizations must now secure the AI agent with the same rigor that's been traditionally reserved for the human endpoint. As AI browsers become deeply integrated into corporate workflows, managing browser-based security risk requires absolute visibility into delegated agent permissions, data boundaries, and real-time behavioral telemetry.

Mitigation framework: How CISOs can reduce enterprise AI risk

The findings described in this SOTI Security report come to a clear conclusion: AI security is no longer about controlling access to a handful of applications. Organizations must govern AI interactions, data flows, identities, extensions, and autonomous agents across an increasingly fragmented AI ecosystem.

Traditional security controls were not designed for AI-native workflows. Effective AI governance requires continuous visibility, contextual understanding of AI interactions, and real-time enforcement that enables employees to use AI safely and does not restrict innovation.

The following framework summarizes the key actions that organizations should prioritize to reduce enterprise AI risk:

- Focus security efforts on AI power users
- Eliminate shadow AI and govern every AI application
- Protect sensitive data at the AI interaction layer
- Secure AI extensions and AI development environments
- Prepare for the era of autonomous AI agents

Focus security efforts on AI power users

Not every employee represents the same level of AI risk. A relatively small population of AI power users generates the majority of AI interactions, shares the most business context, and increasingly relies on AI throughout critical workflows.

Organizations should:

- Identify AI power users across the organization
- Measure how AI is used, not simply which applications employees access
- Monitor prompts, uploads, responses, downloads, and AI interactions involving sensitive information
- Continuously track AI adoption trends and behavioral changes as new AI tools emerge

Rather than applying identical controls to every employee, organizations should prioritize governance around users with the greatest AI exposure.

Eliminate shadow AI and govern every AI application

Enterprise AI now extends far beyond ChatGPT, Copilot, and Gemini. Employees increasingly adopt niche AI applications, AI-enabled SaaS platforms, and personal AI subscriptions without IT involvement.

Organizations should:

- Continuously discover all AI applications in use, including niche and embedded AI services
- Detect shadow AI and personal AI account use
- Require corporate identity federation (single sign-on) wherever possible
- Treat nonfederated corporate logins and personal subscriptions as unmanaged risk
- Perform ongoing vendor risk assessments for AI providers

The goal is comprehensive visibility into every AI service, not just approved platforms.

Protect sensitive data at the AI interaction layer

AI introduces entirely new channels for data movement through prompts, conversations, document uploads, generated responses, and copy/paste activity. These interactions often bypass traditional DLP controls.

Organizations should:

- Inspect prompts, responses, uploads, downloads, and copy/paste activity in real time
- Apply contextual AI-specific DLP rather than relying solely on pattern matching
- Replace “block or allow” approaches with inline AI guardrails
- Apply risk-based segmentation (instead of one universal policy) based on:
 - User risk (AI power users vs. occasional users)
 - Data sensitivity (personally identifiable information, source code, financial information, IP)
 - AI application (enterprise approved vs. consumer AI)
 - Identity (corporate managed vs. personal accounts)
 - Device posture and browser trust

Protecting AI requires understanding what users are sharing and not simply where they are sharing it.

Secure AI extensions and AI development environments

Browser extensions and AI coding assistants increasingly serve as the interface between employees and AI services. Their permissions, credentials, and integrations create an expanding attack surface.

Organizations should:

- Maintain a continuously updated inventory of browser, IDE, and AI extensions
- Evaluate extension permissions, vulnerabilities, reputation, and ownership changes
- Apply adaptive policies to high-risk extensions
- Continuously monitor extension behavior rather than relying solely on static parameters

Extensions should be governed as privileged software rather than as simple browser add-ons.



Prepare for the era of autonomous AI agents

AI agents and AI browsers are rapidly evolving from assistants into autonomous digital workers capable of accessing enterprise systems and executing actions on behalf of users.

Organizations should:

- Inventory the AI agents operating across the enterprise
- Understand what applications, identities, and data each agent can access
- Govern autonomous actions using least-privilege principles
- Monitor agent behavior for abnormal activity and prompt injection attempts
- Extend identity, monitoring, and security controls to AI agents, not just human users

As AI agents become increasingly autonomous, organizations must treat them as a new class of enterprise identity that requires continuous governance.

A practical checklist for CISOs and security leaders

Throughout this SOTI Security report, we've shown how AI is creating new attack surfaces, data flows, and governance challenges. The following checklist consolidates the key recommendations into a practical roadmap to help organizations reduce AI risk while enabling safe and productive AI adoption.

Gain real-time visibility into AI usage

- ☐ Discover all AI applications used across the organization
- ☐ Identify AI usage across browsers, AI browsers, desktop apps, IDEs, and extensions
- ☐ Monitor prompts, responses, uploads, downloads, and AI interactions
- ☐ Track AI adoption trends, power users, and high-risk behaviors

Govern shadow AI and personal accounts

- ☐ Detect unsanctioned and shadow AI activity
- ☐ Maintain approved and restricted AI application lists
- ☐ Identify AI usage through personal identities and unmanaged accounts
- ☐ Detect corporate users who are accessing personal AI subscriptions
- ☐ Establish governance policies for sanctioned and unsanctioned AI services

Prevent sensitive data exposure

- ☐ Monitor prompts for sensitive information before it reaches AI models
- ☐ Inspect file uploads to AI tools for regulated and proprietary data
- ☐ Detect copy/paste activity involving sensitive information
- ☐ Apply AI-specific DLP policies across AI applications and browsers
- ☐ Block high-risk data sharing while allowing legitimate AI use

Secure AI agents, AI browsers, and extensions

- ☐ Inventory AI agents, AI browsers, and AI-powered extensions
- ☐ Monitor what data AI agents can access
- ☐ Govern autonomous actions performed by AI agents
- ☐ Detect risky extension permissions and excessive access requests

Control identity and access risks

- ☐ Enforce single sign-on and multi-factor authentication for AI applications wherever possible
- ☐ Monitor personal-to-corporate account crossover
- ☐ Detect session hijacking and identity abuse attempts
- ☐ Continuously monitor AI access for anomalous behavior

Conclusion

Ultimately, the objective of AI security is not to restrict technology, but to secure its enablement. Robust security frameworks should serve as a business accelerator, providing the visibility, data governance, and real-time telemetry required to operationalize AI safely.

Organizations that establish clear visibility into AI consumption, quantify their data exposure, and identify high-risk user profiles will be best positioned to maximize the economic returns of frontier AI while maintaining resilience.

Credits

Research director

Kimberly Gomez

Writing and editing

Maria Vlasak

Review and subject matter contribution

Eyal Arazi

Pooja Gupta

Promotional materials

Ellen O'Brien

Marketing and publishing

Georgina Morales Hampe

Kimberly Gomez

State of the Internet/Security

[Read back issues](#) and watch for upcoming releases of Akamai's acclaimed State of the Internet/Security reports.

Akamai threat research

[Stay updated](#) with the latest threat intelligence analyses, security reports, and cybersecurity research.

Akamai security research

[Read the Akamai security research blog](#) for a rapid response perspective on today's most important research.

Access data from this report

[View high-quality versions](#) of the graphs and charts shown in this report. These images are free to use and reference, provided that Akamai is duly credited as a source and the Akamai logo is retained.



Akamai Security protects the applications that drive your business at every point of interaction, without compromising performance or customer experience. By leveraging the scale of our global platform and its visibility to threats, we partner with you to prevent, detect, and mitigate threats, so you can build brand trust and deliver on your vision. Learn more about Akamai's cloud computing, security, and content delivery solutions at akamai.com and akamai.com/blog, or follow Akamai Technologies on [X](#), and [LinkedIn](#). Published 08/26.