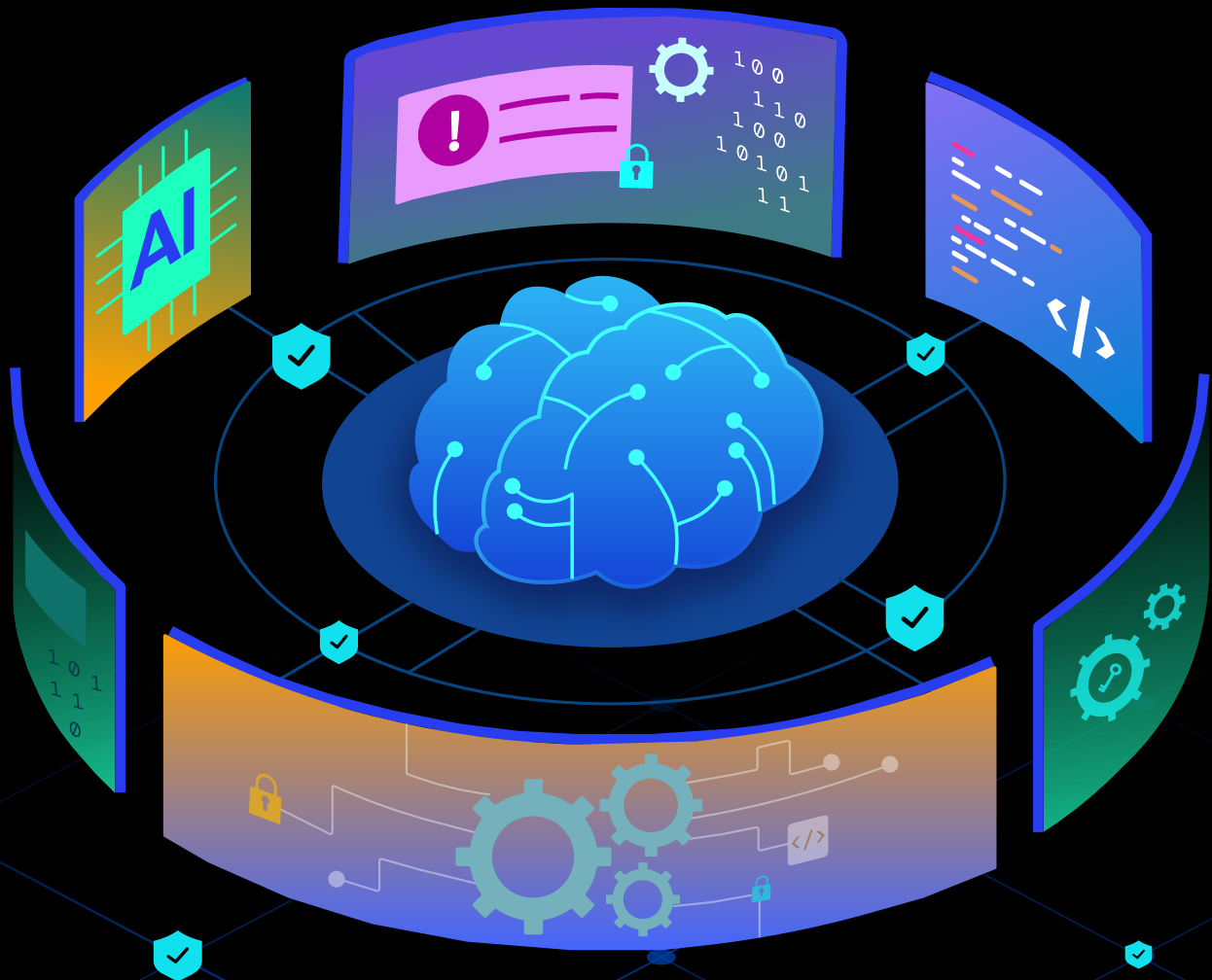


THE STATE OF

AI-Native Application Security 2025



	Executive Summary	02
	Key Findings	03
	1. A Dangerous New Frontier for App Security	05
	2. Solving the Problem of Shadow AI	07
	3. AI App Security: A People or Process Problem?	09
	4. Securing the Future of AI Apps	12
	Conclusion: DevSecOps From Day One	13

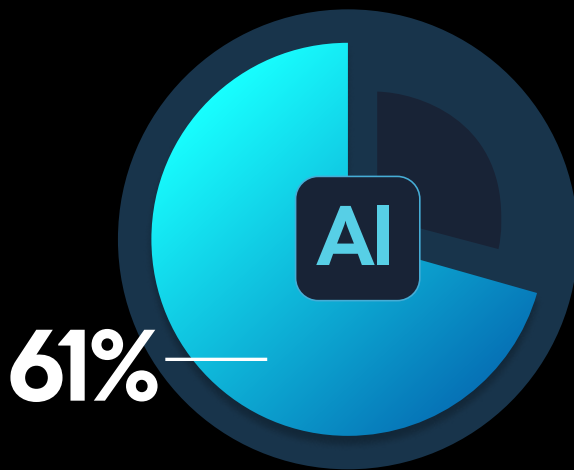
Executive Summary

Enterprises are in the middle of a global AI gold rush. Development teams are scrambling to work Large Language Models (LLMs) and generative AI technologies into their products and workflows at a breakneck pace. However, the opportunity also comes with risk. Our survey of 500 security practitioners and decision-makers across the United States, United Kingdom, France, and Germany finds the rise of AI-native app development has rapidly outpaced enterprise security capabilities.

As AI-native apps flood enterprise environments, security teams cannot keep track of where these technologies are used, how they're implemented, or the vulnerabilities they bring. These blind spots extend across the AI lifecycle, from asset inventory and access controls to API traffic monitoring and threat detection. Organizations are left exposed to an entirely new class of risks they're ill-equipped to defend against.

A breakdown in communication between development and security teams is only exacerbating the problem. Developers often see security as a blocker, bypassing governance and control processes to ship AI-native apps faster. This is causing a proliferation of shadow AI, creating a perfect storm and leaving most enterprises even more exposed to security incidents.

Key Findings

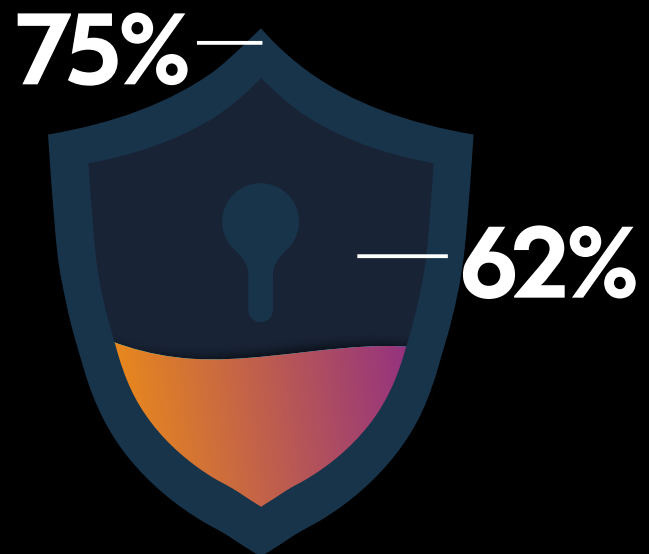


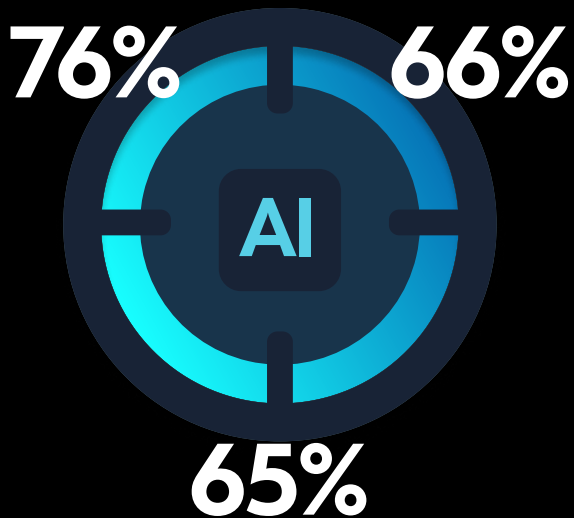
AI-Native Apps Take Over the Enterprise

61% of new enterprise applications are being designed with AI components in mind

The Rise of Shadow AI

- 62% of security practitioners say they have no way to tell where LLMs are in use across their organization
- 75% of respondents say shadow AI will eclipse the security issues caused by shadow IT security as a blocker to AI innovation





AI-Native Apps Are Already Under Threat

- 76% of enterprises have already experienced an LLM prompt injection incident
- 66% have experienced an incident involving vulnerable LLM code
- 65% have experienced LLM jailbreaking

A Chronic Lack of Collaboration

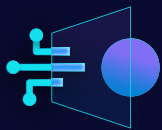
- 43% say developers ensure AI-native apps are always designed with security built in
- 74% say developers see security as a blocker to innovation
- 62% say developers don't have the training to implement comprehensive AI security



1. A Dangerous New Frontier for App Security

In the AI boom, enterprises are increasingly working AI into the foundations of their applications. On average, **61% of new enterprise applications are being designed with AI components in mind**. In a time of such rapid change, there is no standard way for building these apps.

AI Components Being Used in AI-Native Applications



85%

Connecting apps to third-party LLMs, e.g., OpenAI or Anthropic



83%

Connecting apps to third-party AI platforms, e.g., Tensorflow, PyTorch



79%

Integrating apps with custom-built LLMs



78%

Connecting apps to open source LLMs, e.g., Hugging Face

The result is a significantly expanded attack surface that offers threat actors a host of new ways to target enterprises.



Security experts agree, with 82% saying AI-native applications are the new frontier for cybercriminals.



Additionally, 63% believe AI-native applications are more vulnerable to security threats than traditional IT applications.

In this new frontier for application security, risk exposure is skyrocketing. In fact, most enterprises have already experienced incidents where AI-native apps have been targeted.

Attackers Are Targeting LLMs

Most enterprises have already experienced security incidents.

76%

Prompt injection: manipulating LLM prompts to insert malicious code or extract sensitive information

66%

Vulnerable LLM code (or vulnerable third-party code used by the LLM): which can be exploited by attackers to execute malicious code

66%

Unbounded consumption: where an LLM allows users to conduct excessive and uncontrolled inferences, leading to denial of service (DoS), economic losses, and service degradation

65%

Shadow AI: where AI is connected to data and systems without security teams' knowledge

65%

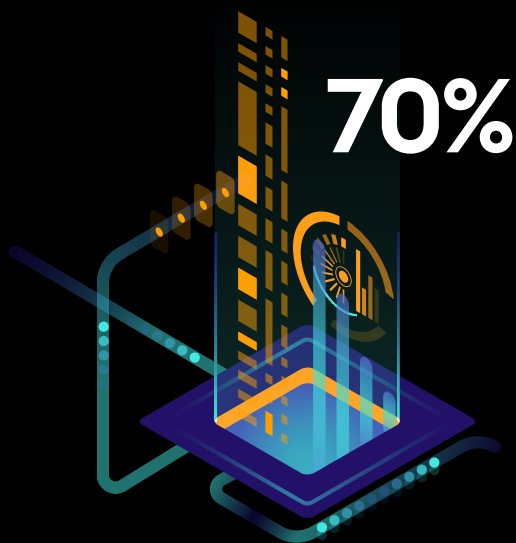
Jailbreaking: manipulating LLM prompts and causing the AI tool to disregard safety protocols altogether

63%

System prompt leakage: where an LLM voluntarily gives up sensitive information

2. Solving the Problem of Shadow AI

Enterprises are struggling to gain visibility into rapidly increasing AI use. As these technologies plug into various enterprise systems and access more data, a complex web of connections is making it more difficult to spot the security and compliance gaps.



70% of respondents say it seems like a new API connects an LLM to sensitive data every day in their organization. This is creating a shadow IT problem at scale.

75%

say shadow AI will eclipse the security issues caused by shadow IT

74%

say AI sprawl will blow API sprawl out of the water when it comes to security risk

72%

say shadow AI is a gaping chasm in their security posture

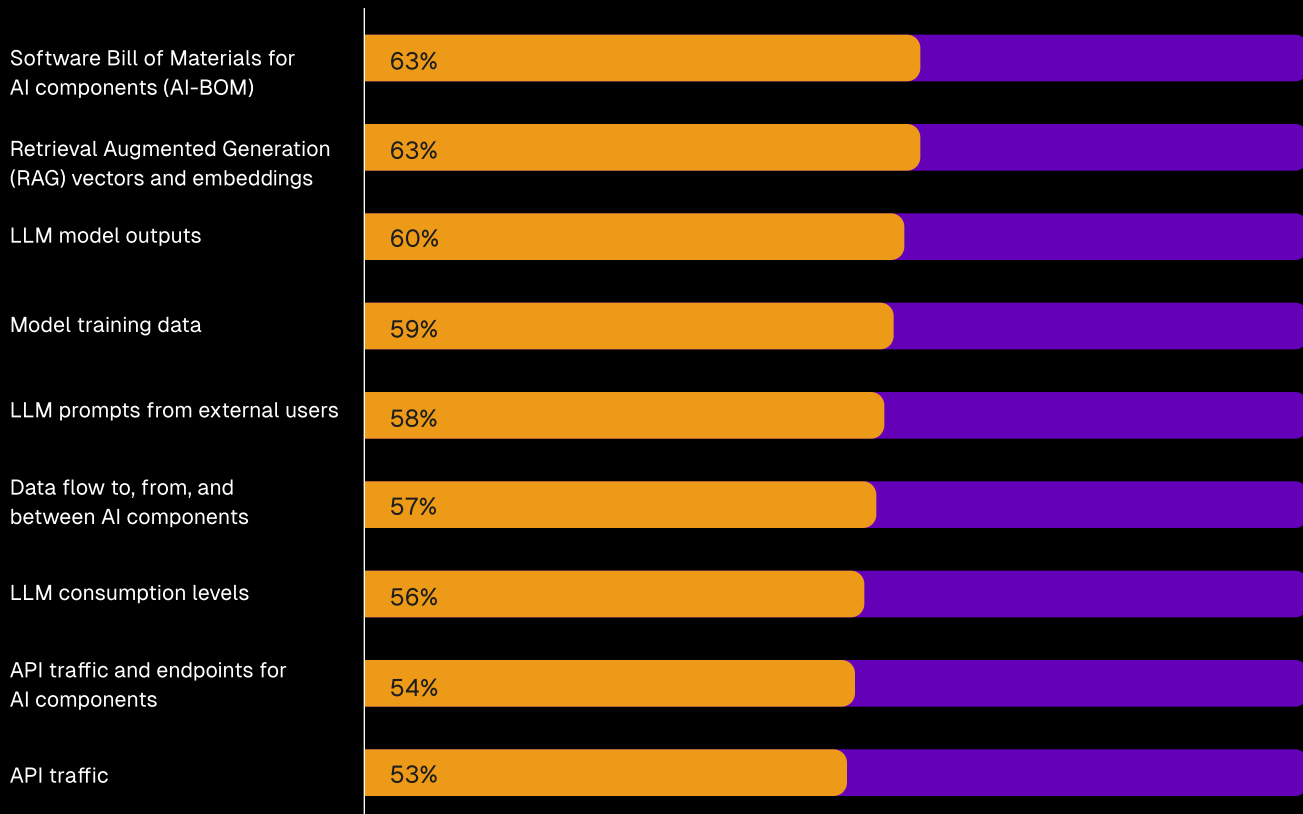
66%

say they are flying blind when it comes to securing AI-native apps

62%

say they have no way to tell where LLMs are in use across their organization

Despite acknowledging the risks, many organizations' security teams do not have full real-time visibility into the most critical aspects of AI-native applications.



Without real-time insight into the behavior of their AI components and the APIs that connect them, enterprises are left exposed to emerging threats like LLM jailbreaking, sensitive data leakage, and AI Denial of Service attacks.

3. AI App Security: A People or Process Problem?

To regain control over the security of their applications, enterprises need to drastically increase real-time visibility into their data, APIs, and AI components. At the same time, people and processes must adapt to help enterprises secure AI-native apps.

People: Security Teams Racing to Keep Pace with AI

To defend against AI threats, teams require a higher grade of visibility and control over cloud and API security than ever before. But, they also have new skills to learn and processes to implement.



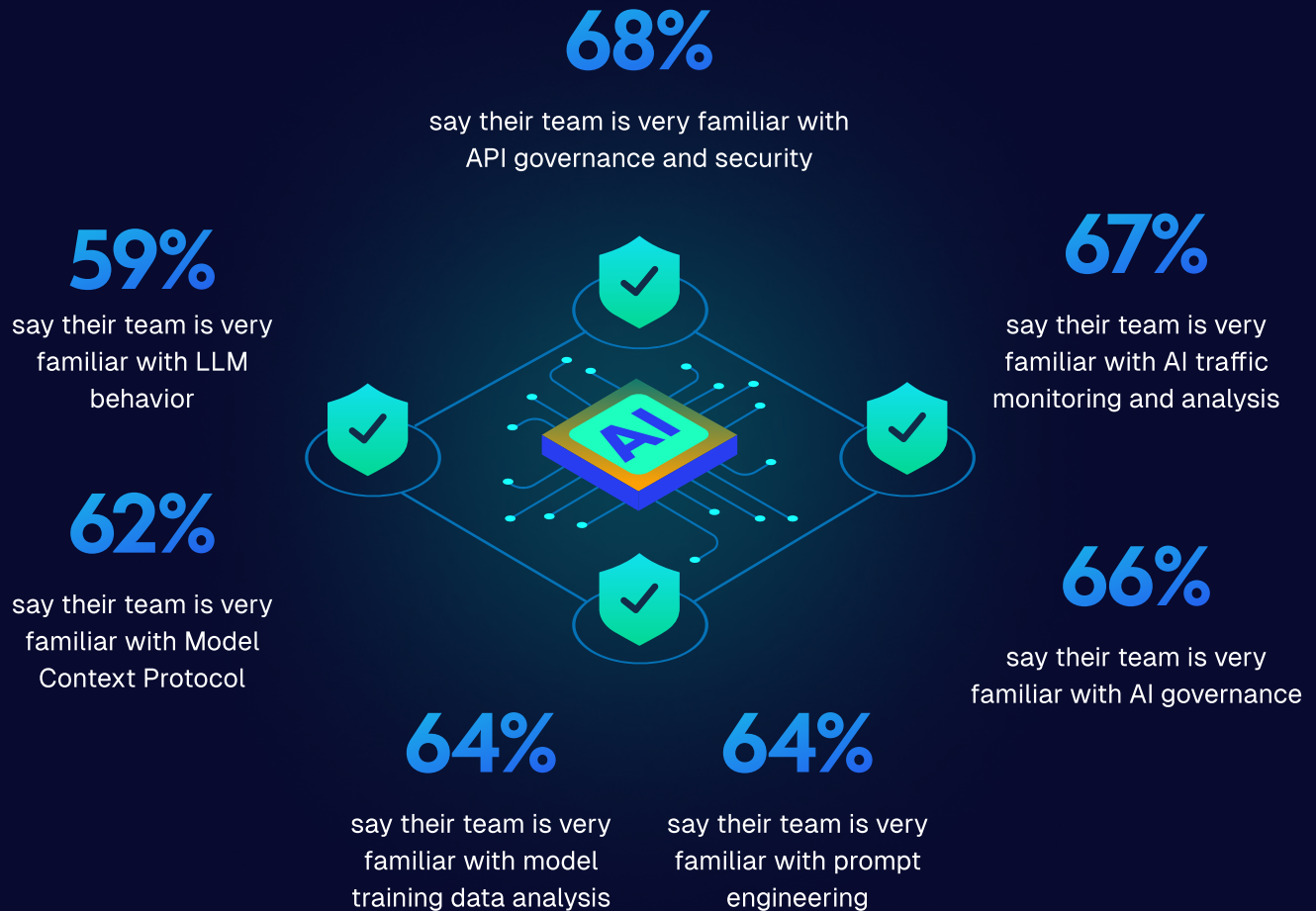
75%

of respondents say AI-native apps advance so quickly that security teams are always on the back foot

75%

say security threats for AI-native apps are a whole new kettle of fish, as they never had to think about prompt injection for traditional apps

In response, security practitioners have been brushing up on their skillsets to adapt to the new normal, and are quietly confident in their knowledge around AI-app security:



Processes: A Lack of Collaboration With DevOps

While security teams work to keep up with AI advancements, they are also struggling to win a battle of hearts and minds with development teams, where governance processes and controls are often seen as a hindrance to progress.



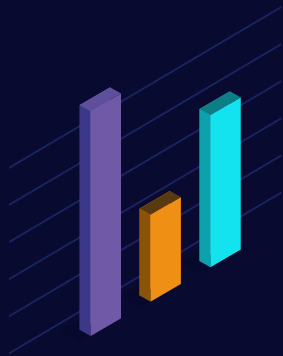
In the rush to deploy AI, a serious disconnect is forming between security and development teams. Currently, security teams feel they are out of the loop, with 62% of respondents saying their developers aren't taking responsibility for securing AI-native applications.

This communication breakdown is apparent from the start of the process of building AI-native applications.

43% of organizations say developers ensure AI-native apps are always built with DevSecOps principles in mind (i.e., with security built in).



When creating a new application,



34%

just over a third (34%) of developers let security teams know before they get started

53%

will notify security teams before going into production

14%

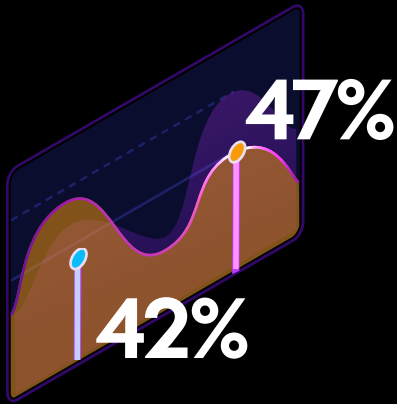
will only inform security teams after the app has gone into production, or when a security incident has occurred

There is also acknowledgement of a skills gap. 62% say their developers don't have the time and 62% say their developers don't have the training to understand AI-native application security.



With developers not looping security in at the start of their projects, there are widening gaps for malicious actors to exploit in today's AI-native applications.

4. Securing the Future of AI Apps



Against a rapidly shifting AI landscape, the industry is working to catch up quickly. **47%** of security teams say regional AI regulations will be highly effective at enforcing secure application development practices, and a further **42%** suggesting they will be moderately effective.

If regulators are doing enough, that leaves the ball in the enterprises' court to ensure they can meet emerging compliance standards and plug any gaps that arise in their AI-native application security posture. To succeed, enterprises are focusing on ways to improve visibility, bolster resilience, and better protect AI-native apps.

The most important step is to identify where LLMs are used across the organization. Currently, security teams use the following methods to achieve this:

63%

are monitoring
access controls
for AI agents

59%

are monitoring
API traffic

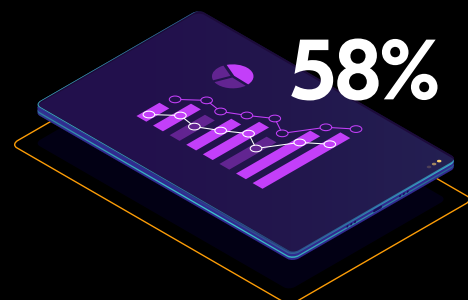
57%

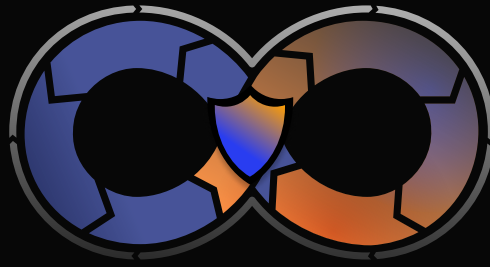
are carrying out
inventory checks
with developers

47%

are checking with
Finance to track
AI spend

Most enterprises have also put measures in place to evaluate and manage the posture of AI-native apps, with **58%** establishing governance policies. These will be critical for securing future applications as they are built.





Conclusion: DevSecOps From Day One

While AI progress is moving faster than the human eye can follow, it is leaving critical security gaps in AI-native applications that need to be plugged fast.

To reduce the risk of shadow AI and the impact of related incidents, enterprise security and development teams must work together to boost visibility into AI components and implement DevSecOps processes from day one. This means:



Ensuring that security is built into AI-native apps from the start, with clear governance policies and communication between developers and security.



Discovering all new AI components as they appear and ensuring they are monitored and logged.



Achieving real-time visibility into AI components and the services they communicate with, focusing especially on API traffic.



Carrying out dynamic application security testing (DAST) to identify security risks prior to production.



Protecting AI-native apps in production, inspecting prompts and monitoring responses to reduce sensitive data disclosure.



Methodology

This report is based on a survey of 500 security practitioners and decision-makers responsible for securing AI-native applications, commissioned by Harness and conducted by independent research firm Sapio Research. The sample consists of 200 respondents in the United States, and 100 each in the UK, Germany, and France.