

Latio

AI Security Market Report

Reporting Period
Q2 • 2025

Prepared by
James Berthoty

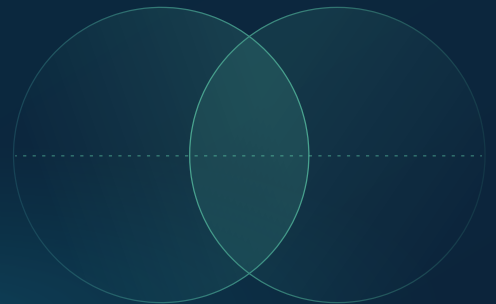


TABLE OF CONTENTS

REPORT INTRODUCTION	3
AI USE CASES	4
AI ADOPTION & SECURITY NEEDS	11
OLD SOLUTIONS, NEW PROBLEMS	15
MARKET MAP AND LEADERS	18
TOOL BUYING FLOW CHART	29
VENDOR BREAKDOWNS	34
CONCLUSION	45

AI Security Landscape

AI security in 2025 has been defined by a disjointed landscape of solutions to unclear problems.

A venture capital fueled marketing frenzy has led to widespread confusion across social media, with significant misunderstandings around both use cases and best practices.

This report cuts through the noise to clarify what's been lost in the murkiness of so-called "AI-TRiSM."

While the AI Trust, Risk, and Security Management category was created to bring a range of solutions under a single label, it has often flattened the important distinctions between tools, especially when it comes to what problems vendors are solving and how.

Most of AI Security might feel new, but many of the underlying challenges are quite familiar. In nearly every category we'll cover, there are existing tools that offer similar functionalities to various startups.

While that may suggest security teams should wait to adopt new solutions, it will take time for these traditional vendors to match the pace and specialization of newer, AI-native offerings.

As with any security decision, the right choice depends heavily on three important factors:

- **Your specific risk profile**
- **Technology stack**
- **Organizational priorities**

By the end of this report, you'll walk away with two things:

- * *A clear understanding of when and why to use an AI security tool, including what specific use cases these tools are designed to address.*
- * *A structured overview of the vendors currently in the space, what they offer, and how to evaluate them based on your needs.*

We'll also include a simple decision flowchart to help guide your tool selection based on real-world scenarios.

We hope this information is helpful, and we thank you for using Latio as your source for trusted industry insights.

AI SECURITY USE CASES

Let's start by outlining the different AI security use cases. These can be separated into four major categories:

01

End User Data Control

1. Data Loss Prevention
2. SaaS Access control
3. Secure Code Creation

(IT teams)

02

AI Posture Management

1. Infrastructure Discovery
2. ML-BOM/AI-BOM
3. Data Pipeline Posture
4. Static Code Testing

(Infrastructure teams)

03

Application Runtime

1. Prompt Injection Protection
2. Visibility into runtime models
3. Authn/Authz
4. Dynamic Testing

(AppSec/DevSecOps)

04

AI for Security

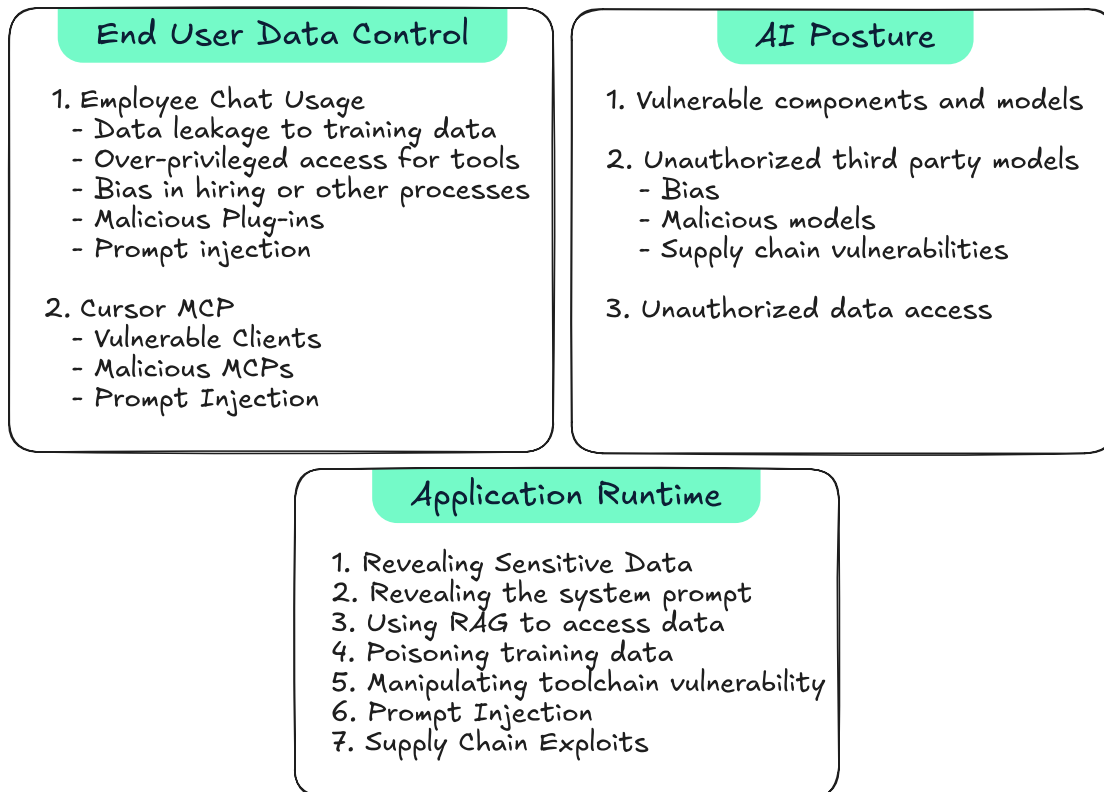
1. AI for SOC
2. AI for Vulnerability Management
3. AI for AppSec

(SOC and AppSec)

As a brief aside, you could argue for a fifth category, **policy and compliance management** around AI concerns, but we're choosing to exclude it from this report. It's a smaller use case and typically less relevant to security engineers.

Risk Types

Top AI Attacks per Category



Each category of AI tooling is designed to protect against specific risks for these categories. On the end user data side, tools protect employee endpoints against manipulation, or stop employees from sharing unauthorized materials with unauthorized AI systems.

On the posture side, security teams struggle to get deep insights into how models are being developed and deployed, and what datasets these models have access to. While many teams are relying on third party models, self hosted models are especially vulnerable to different poisoning or supply chain attacks. Additionally, models can have vulnerabilities the same as any other code packages.

Finally, runtime application security is the most at risk for real world attack, especially once a system is wired up to internal data. Early iterations of AI applications were low risk, as they merely surfaced a user's data back to them; however, agentic architectures have rapidly increased risk as agents take actions on behalf of users, and have access to sensitive data.

End-User Data Control

Gaining Visibility into Employee AI usage is where most enterprises begin their AI security journey. The natural starting point is monitoring and controlling the input of sensitive data into third-party models. Early in the adoption of AI, it was unclear whether models hosted on platforms like Hugging Face would become the standard or if proprietary options like ChatGPT and Claude would dominate. Over time, it has become clear that most employees are gravitating towards a smaller number of options - usually their workforce AI tool, such as Microsoft Co-Pilot or Gemini, or ChatGPT.

The core security concerns in this category include:

- * Monitoring for sensitive data being shared with third-party tools
- * Ensuring those tools are configured not to use sensitive data for training
- * Detecting suspicious behavior, such as a user falling victim to a prompt injection attack

As specific AI tools have matured, their use cases have also evolved. For example, Cursor is widely used for code generation, while Microsoft Copilot has gained traction for enterprise no-code AI applications.

For tools like Microsoft Copilot, the main concern is **permissions management**. Users may unknowingly have access to sensitive data or create applications with broader permissions than intended, which the AI can then leverage to fetch privileged data. This has highlighted a growing need for tighter **SaaS identity and access management**.

For **secure code generation**, typically through Cursor, many organizations have accepted the productivity benefits despite the security tradeoffs of exposing code to third-party models. The focus has shifted towards **maintaining code quality** and ensuring AI-generated code aligns with organizational standards.

To summarize, end-user data control revolves around three key subcategories:

1. **Data loss prevention** via AI chatbots
2. **SaaS access control** for AI tooling within productivity platforms
3. **Secure code creation**, particularly with developer-facing tools like Cursor

AI Posture Management

This second category focuses on securing the infrastructure that supports AI itself. Much like the first, it has evolved rapidly as companies have begun to standardize deployment patterns. In the early stages, organizations used a mix of local Hugging Face models, open-source alternatives, and cloud-hosted models from providers like OpenAI and Anthropic. Over time, API-based models have become more the norm, driven largely by performance advantages. The main exception involves teams hosting local models for cost savings, and utilizing frameworks like LangChain for creating agentic architectures.

CNAPP vendors like Wiz were early to recognize the opportunity of this space, coining the term **AI Security Posture Management (AI-SPM)**. At first, it referred to discovery, helping organizations identify what models were in use and where they were deployed. But it has since expanded to include lifecycle management of AI development environments, including monitoring Jupyter notebooks, Airtable pipelines, and other complex workflows beyond just EC2 instances.

The risks at this layer are mainly about **misconfigurations** and **unauthorized access** to models. A notable example is the [PyTorch vulnerability](#) discovered by Oligo Security, which demonstrated how misconfigured AI services could become actively exploited threats.

Another emerging concern is **model poisoning** and supply chain integrity. This has led to standardization around "ML BOMs", machine learning bills of materials, used interchangeably with "AI BOMS," to help track and authorize model provenance. Snyk and other ASPM vendors like Cyscale and Apiiro have introduced capabilities to detect and map the risk of using AI models, such as determining their dependencies and datasets. The goal is to prevent tampering or unintentional biases, such as tweaking a model to favor a specific product in a shopping platform.

In short, AI posture tools aim to secure the **infrastructure and lifecycle of AI systems**, helping organizations identify, manage, and harden what powers their AI.

Application Runtime Protection

Runtime Application Protection is most relevant for companies building and deploying their own AI-powered applications. Initially, many of these applications were basic chat interfaces, offering minimal security risk. Users could ask questions and receive answers, but the AI wasn't granted access to any sensitive systems or data, only the prompt itself.

Today's applications, however, are significantly more capable. Modern architectures allow AI agents to access internal systems, retrieve sensitive user data, and even execute code. This shift has **dramatically increased the risk profile** of these tools, prompting a wave of concern among security leaders.

The most common attack type here is **prompt injection**, where the model is tricked into performing actions outside the original intent. However, there are other threats as well, including **model reconnaissance** and **stealthy bias injection**, which can be harder to detect.

In response, the industry is focused on two key areas:

- * *Monitoring for malicious activity at runtime*
- * *Testing AI models pre-deployment to identify vulnerabilities*

This has given rise to a niche group of vendors performing security testing for AI models, similar to how DAST tools scan web applications. These tools probe models for bias, robustness, and attack resistance before deployment.

Shockingly, this has also revived a lightweight form of **RASP-style** protection for AI applications, with runtime monitoring agents guarding against real-time threats.

AI for Security

While most security tools are developing their AI story for investors, we also wanted to highlight both startups who are building from an AI-native approach, and incumbents that are leading in AI adoption.

AI is strongest when it comes to summarizing large volumes of data for human consumption. Therefore, most innovation in AI for security has focused on time-consuming analytical tasks. The two main areas where this has taken hold are **security operations centers (SOCs)** for incident analysis, and **application security** for code analysis.

AI performs well in both cases when it has access to all relevant data. In fact, it often produces summaries that surpass human capabilities, especially since these tasks are frequently handed to analysts with limited knowledge of the full range of technologies, platforms, and tools involved.

As a result, there has been a surge in development around **AI-driven code fixing tools** and **SAST solutions**. While we covered AI autofix tools in [a previous report](#), this section focuses on the emerging category of **AI-native code scanning**.

SAST solutions built from the ground up with AI are particularly well-suited to uncovering authentication-related vulnerabilities and other context-aware issues that traditional scanners often miss.

Given that a majority of vulnerabilities, especially those in the OWASP Top 10, are related to authentication, and that traditional static analysis struggles to detect them, we believe AI agents will play a central role in the future of application security scanning.

In the **security operations** space, AI-powered analysis tools often appear highly advanced in demo environments. However, in real-world use, results have been more mixed. One of the most common challenges in a SOC is the discovery of **missing data** during an investigation. In these cases, AI can be prone to hallucinating conclusions or generating analysis that lacks actionable value, particularly when it doesn't have access to key logs or unfamiliar data sources.

That's why we prefer solutions where all data is already centralized, allowing the AI to query everything it needs to construct an accurate investigation timeline. Two notable vendors here are **Exaforce** and **AI Strike**. Both offer robust data platforms that embed AI throughout the entire SOC workflow, from analysis and detection engineering to SIEM health and optimization.

A final area worth highlighting is **vulnerability management**. Several startups have focused on using AI to better organize or enrich existing data.

Two good examples are **Phoenix Security** and **Opus Security** (acquired by Orca Security). Phoenix uses AI to map CWEs to vulnerabilities, improving prioritization logic. Opus applied AI to power advanced workflows for data correlation and vulnerability remediation. A newer entrant worth mentioning is **MazeHQ**, which takes an end-to-end AI approach to vulnerability analysis.

Now that we've laid out the core use cases of AI security, let's look at what the priority is for industry leader adoption.



AI ADOPTION & SECURITY NEEDS

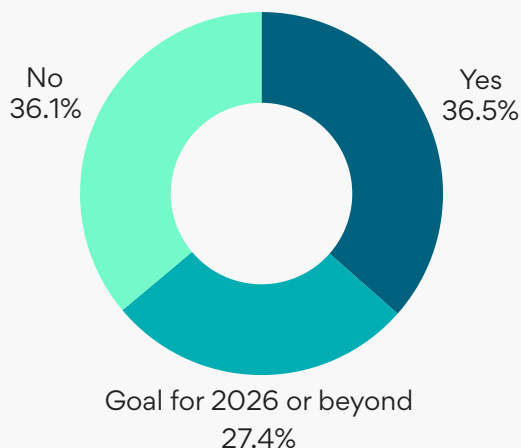
With use cases defined, we polled security leaders to understand how AI security is being prioritized within their organizations. The results reveal an emerging focus on AI risks, particularly around application development.

Audience Reports Pt. 1

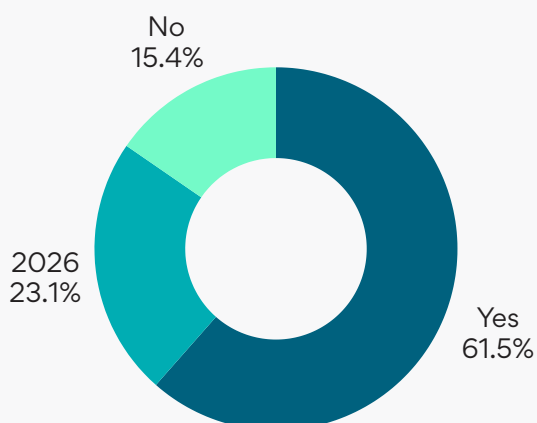
Strong Exploratory Interest

First, with securing AI workforce usage, security leaders were split almost evenly between prioritizing it this year, next year, and not having any concerns at all. Prioritizing workforce security was closely tied to industry and security maturity. **Teams with a high level of existing endpoint security and DLP were much more likely to prioritize AI, while those without were unlikely to prioritize it only for AI usage.**

Is Securing AI workforce usage a 2025 goal?



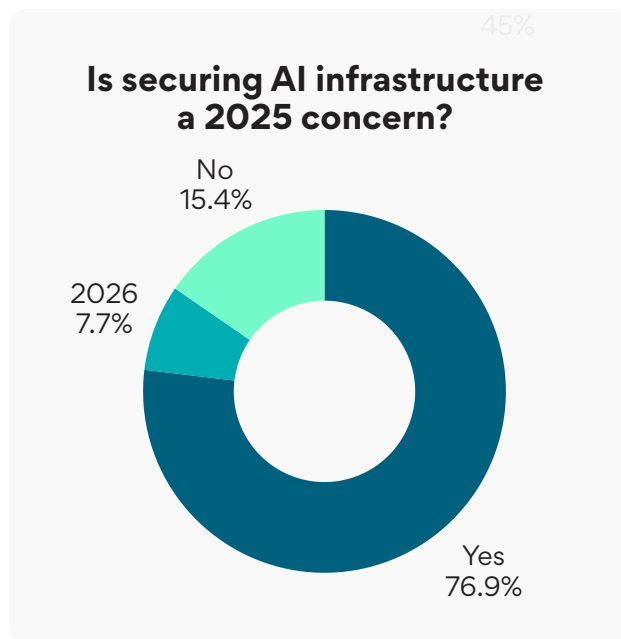
Is Securing AI Application Development a 2025 goal?



Second, a majority of teams had securing AI Application Development as a 2025 goal. It's clear that practitioners are quickly building guardrails for AI application development and best practices. The rush to this emerging category has led to the quick deployment of MITRE and OWASP guidance for AI application development.

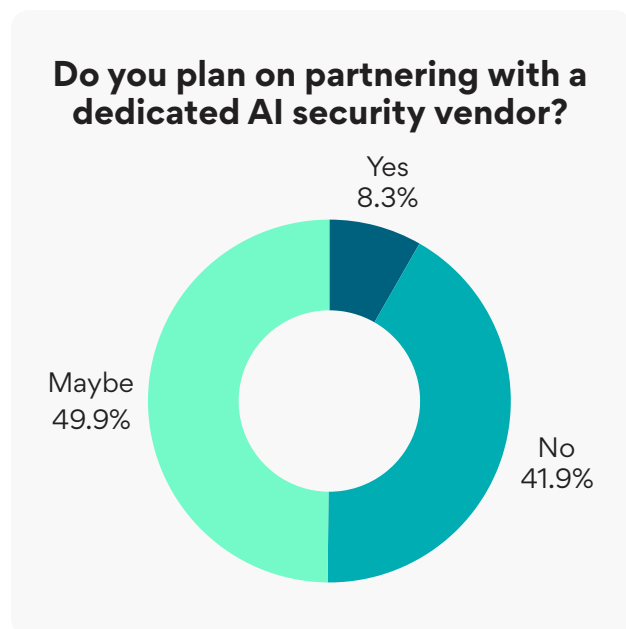
Audience Reports Pt. 2

Even more than AI Application Development, securing the underlying infrastructure was the primary goal of security teams this year. Teams are seeking to fully understand their AI infrastructure and correct any misconfigurations, much like the early days of cloud security.



Immediate Buying Needs

While the majority of respondents stated that AI Security was a short-term goal, they did not prioritize it as the most urgent goal, nor did they unanimously want a vendor to help them accomplish it at this point. It's our interpretation that this is due to the AI security problems still being made clear - it's unclear to practitioners what they even want a solution to look like yet.



When asked about plans to partner with dedicated AI security vendors, the most common response was “maybe.” While few respondents said they were actively pursuing partnerships, the high number of tentative responses suggests that interest is strong, but commitment is still forming. Many teams are waiting for the AI Security problem to become more defined before looking for any partnerships. **They're sure there will be an AI security problem, but aren't sure what that will be.**

Audience Reports Pt. 3

The data shows that security leaders are taking AI seriously, especially as it becomes more deeply embedded into internal products and workflows, but many are still navigating what responsible adoption and protection should look like. **Less than 10% of respondents had AI Security as their most immediate concern, but almost 50% thought it would be soon.**

Is AI your most immediate security concern?

■ Not yet, but soon ■ No

46%

40%

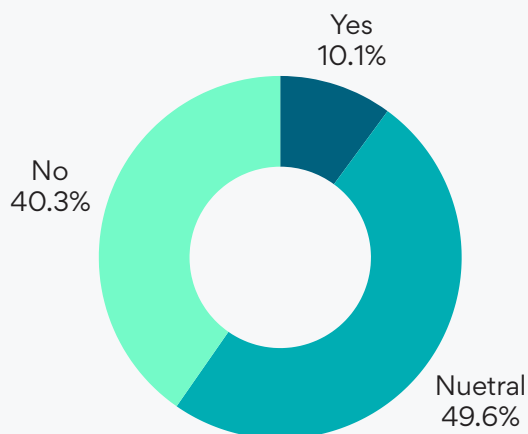
What dedicated resources will you allocate to achieve AI focus?

* Internal Employees

* Evaluating VARs

* Tool Evaluation

Do you care if a vendor uses AI?



Finally, security users largely don't care about if their vendor is using AI or not. Security teams care about outcomes, not the process by which it is achieved!

OLD SOLUTIONS, NEW PROBLEMS

While it's tempting to look at emerging threats and feel the need to immediately invest in new tools, the reality is a bit more nuanced. Many of the detection capabilities we now associate with AI security actually map to existing, fairly mature technologies. While incumbent vendors in these spaces may not offer the same feature depth as newer, AI-specific players, it is a viable strategy for security leaders to wait for these vendors to catch up rather than rushing to adopt a new stack.

Historical Categories Mapping to AI Use Cases

End User Data Control <i>Incumbent Categories</i> 1. Firewalls 2. ZTNA 3. DLP 4. SaaS Security 5. IDE Plugins 6. Email Security 7. Secure Browser	AI Posture (Infrastructure) <i>Incumbent Categories</i> 1. CSPM 2. SCA 3. SAST 4. ASPM 5. CTEM 6. Firewalls 7. WAF
AI for Security <i>Incumbent Categories</i> 1. SIEM 2. CDR/ADR/CADR 3. CTEM 4. ASPM/SCA/SAST	Application Runtime <i>Incumbent Categories</i> 1. ADR 2. CADR 3. WAF 4. DAST 5. API Security 6. RASP

Old Solutions, New Problems

Starting with end-user data control, browser monitoring has been a security capability for a long time. Whether through network traffic inspection at the VPN level, secure browsers, IDE plugins, or SaaS integrations via APIs, **many existing tools already have access to the key underlying information: who has access to what data, and what data employees are sending to external services.**

Firewall vendors, for instance, have quickly adapted to basic use cases like **monitoring employee AI chatbot usage and enforcing access policies for approved models or endpoints.** Similarly, vendors focused on SaaS security or DLP already offer solid detection for permissioning issues and sensitive data exposure. That said, monitoring tools like VS Code plugins remain a weak point. Developer workstations often lack the protections applied to other endpoints, and keeping tabs on AI-powered coding tools is still an open challenge for most enterprises.

Next, when it comes to **AI posture management**, traditional security vendors have a significant advantage. **Many already integrate with cloud environments and scan workloads, making it a logical step to start detecting AI-specific signals like package usage, model deployment, or exposed APIs.** One area where traditional tools fall short, however, is the broader data security landscape that newer vendors like Noma are

targeting, covering everything from model deployment to data engineering and lifecycle management.

On a similar note, **SCA scanners** already provide visibility into package usage and dependencies. They're well positioned to identify which models are being used and where they are deployed within the codebase.

Turning to **AI application protection**, it's important to recognize that **most AI apps are just collections of SDKs and API calls.** This means traditional application security solutions are already moving into the AI space, building runtime protection into their existing offerings. Still, since runtime application protection is relatively new, nearly everyone is still in catch-up mode.

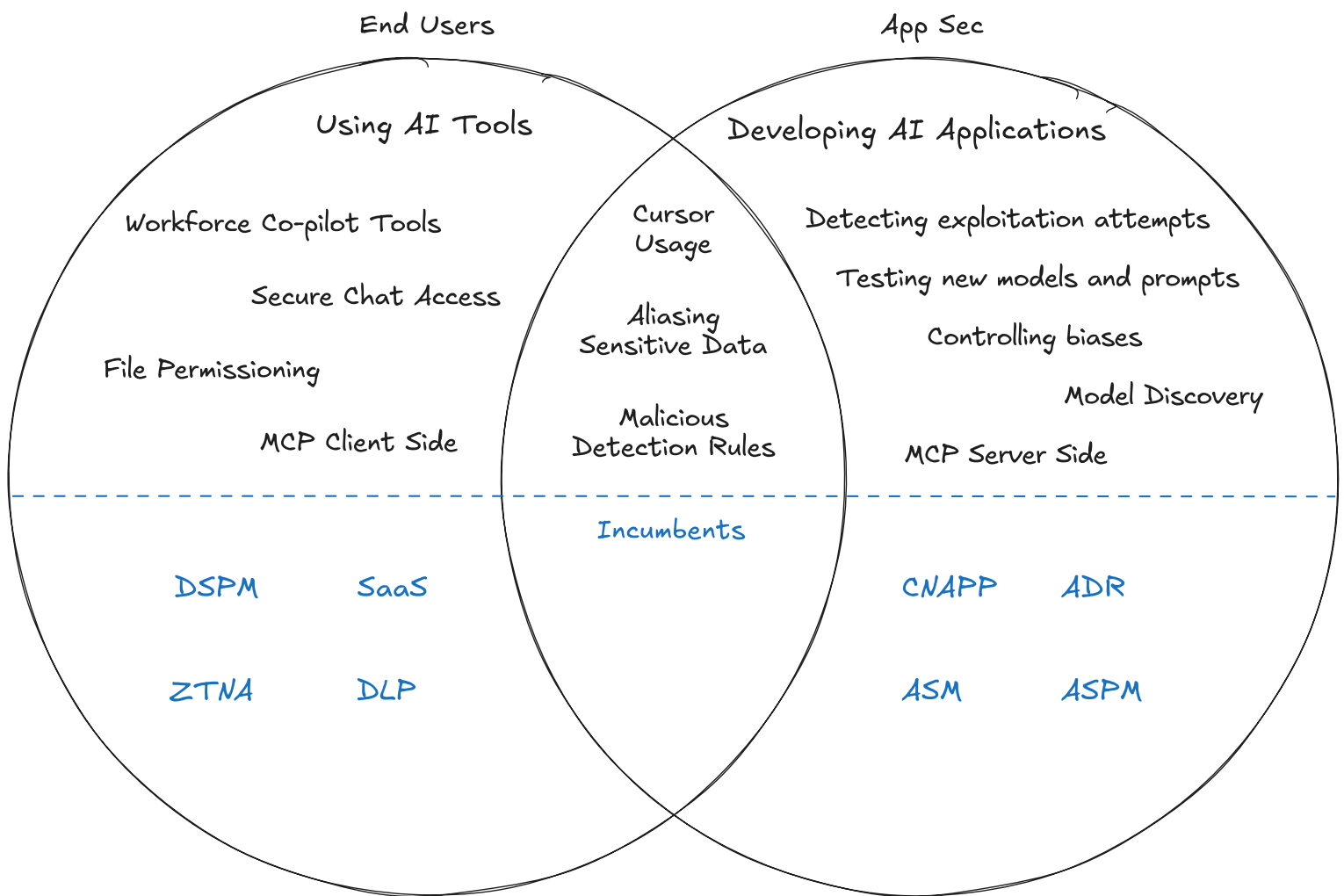
At a baseline, **WAFs will evolve to include detections for things like prompt injection**, though creating effective rules will be challenging. These attacks are highly context-specific and difficult to detect using traditional signature-based methods.

Application Detection and Response (ADR) solutions are particularly well-positioned to handle runtime protection for AI applications, but dedicated AI runtime tools are currently ahead in terms of coverage and detection accuracy.

In particular, **AI security testing** is gaining traction as enterprises seek to understand and proactively identify how these attacks actually occur.

Finally, both SAST and SIEM vendors are implementing their own AI-assisted solutions, though it remains to be seen whether they can match the speed and effectiveness of tools built natively with AI at their core.

Use Cases



Latio



MARKET MAP & LEADERS

Market Map & Leaders

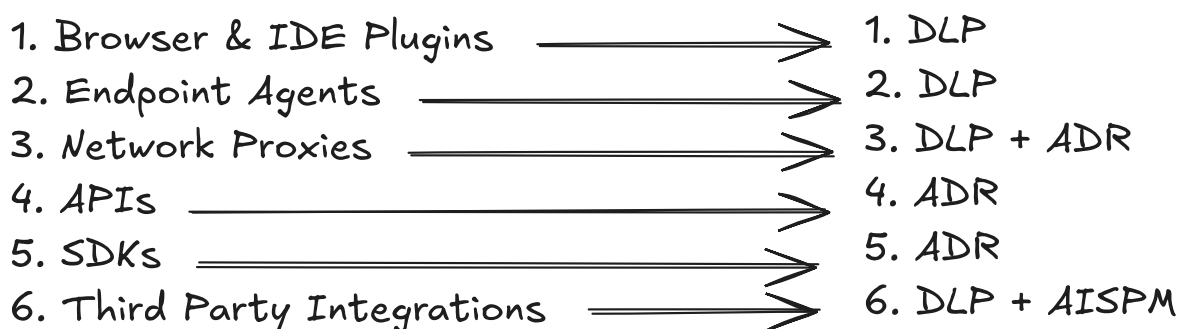


All diagrams are [available here](#).

Implementation Methods

Before diving into specific platforms and tools, it's important to understand the **different implementation methods** vendors use to achieve their goals. The most common approach for enterprise visibility and data loss prevention is through **browser or IDE plugins** that monitor employee usage of third-party chat systems. For application protection, an SDK often implements visibility and protection around the AI usage. **Most unified tools don't quite offer a full SDK, as their offering in this space is newer only offering a basic API.**

Implementation Methods



Many vendors also provide a **network proxy**, which routes user requests to approved AI models. This proxy serves as a front end for monitoring model usage and detecting risky prompts. In many cases, this architecture also enables **runtime application protection**, since network requests can be intercepted and inspected in real time.

For **runtime application protection**, most providers offer APIs as well as SDKs, which function as custom wrappers around those APIs, that developers can integrate into their applications. These SDKs help monitor for prompt injection and other attack types, while also enabling testing functionality before deploying AI capabilities in production.

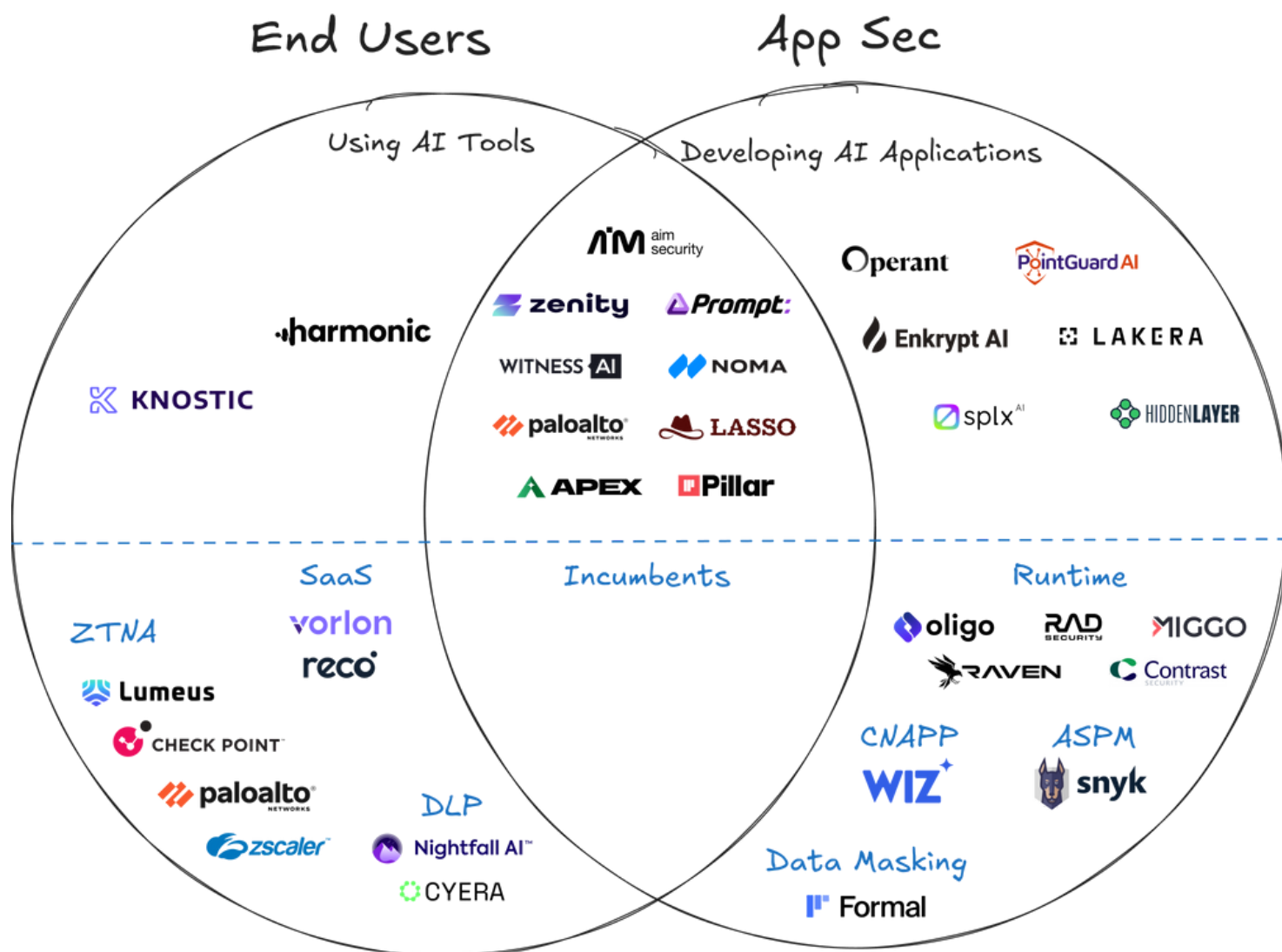
Finally, **third-party integrations** are a common feature for both data loss prevention and posture management.

On the DLP side, these integrations often focus on tools like **Microsoft Copilot**, enabling visibility and control over user-generated AI applications. On the posture side, vendors typically integrate with **cloud environments** like AWS to identify which models are in use and how they're deployed. Some tools also connect with **source code management platforms** like GitHub to monitor which AI models are embedded in the codebase and where.

Trade-offs

Like most other categories in security, **there are trade-offs to using platform providers compared to dedicated ones**. Most platform providers benefit from having a unified detection engine across end-user detections and application detections. Additionally, they benefit from being seen as a single AI Security partner for large enterprises and focus on accelerating secure AI adoption.

Of the unified platform tools, almost all started as end-user protection oriented and have since shifted to runtime application protection. This is a sign of where the market is headed, namely towards application protection. Additionally, most unified providers are much less mature at application protection than the dedicated providers.



The Unified Approach



All-in-one providers tend to offer capabilities across multiple implementation areas, though their maturity often varies from one function to another. Most providers in this category started their development in the end-user protection category by offering browser plugins for monitoring employee AI usage. Since then, they've aimed to evolve into larger use cases - from MCP protection to application protection via API integrations.

For example, **Prompt Security** and **Lasso** both provide **browser plugins** and **endpoint agents** to monitor employee AI usage, along with a **network proxy** to redirect and inspect AI traffic. However, unlike dedicated **application runtime security** vendors, they don't support discovery of AI through code integrations and offer only minimal application protection capabilities.

Palo Alto Networks recently entered into this unification category through their acquisition of Protect AI, who was focused primarily on model testing. Through this acquisition, Palo Alto provides AI security through both their network security offerings as well as an option for application security adopters.

The value of having a unified tool will depend on your use case. In general, these tools excel for companies that were first concerned with employee monitoring and sensitive data detection use cases. Their offerings are most mature at detecting and blocking suspicious or unwanted AI activity. They have strong detection engines that extend into developer workflow options via SDK and API usage.

Broadly speaking, these platforms aim to serve as **long-term partners** for enterprises managing AI security holistically.

End User Data Control

Vendors focused on **monitoring end-user data** typically offer multiple ways to surface insights into risky or inappropriate AI usage. These use cases span **data loss prevention (DLP)**, **insider threat detection**, and identifying **user-facing prompt injection attacks**.

Three of the vendors highlighted below exemplify different approaches to the problem: Knostic, Zenity, and Harmonic. **For clarity, most of the unified providers take a similar approach to Harmonic.**

Harmonic concentrates on monitoring direct usage of AI by employees, whether in chatbots, code assistants, or other AI tools. It monitors AI interactions via **browser extensions**, and other mechanisms that can provide comprehensive visibility into AI interactions. They classify types of security data that shouldn't be shared to different providers, and then alert or block the user from sending that data.

End User Data Control



Zenity, originated as a **no-code security platform** and has been particularly well-positioned to secure **internally developed AI applications**, such as those created through **Microsoft Copilot**. Zenity monitors both the configuration and runtime behavior of these applications, making it a strong fit for organizations enabling internal AI development without introducing heavy onboarding friction.

Knostic focuses on contextual data access for AI usage. While they don't directly block or stop any access attempts, they build a permissions model based on use cases, making deploying AI security solutions at scale much more doable. For example, making sure only HR employees can ask questions about compensation, rather than blocking it holistically.



AI Posture Management

AI Posture (Infrastructure)



Wiz brilliantly coined the term “**AI-SPM**” early in the cycle of AI adoption, capitalizing on the wave of model releases to establish an early brand around discovering and monitoring AI usage across the enterprise. Their **agentless scanning approach** was well-suited to the kind of **basic visibility** security teams needed: namely, understanding **what AI technologies are in use and where**. Wiz’s tooling can detect where models are deployed and help security teams gain broad visibility across cloud workloads. 

A smaller, niche category that has emerged here is **ML-BOM (Machine Learning Bill of Materials)**, also called AI-BOM, tooling. This concept is a natural extension of traditional **SCA (Software Composition Analysis)**, focused on tracking model provenance. Since many models are either **imported as code libraries** or **installed directly on systems**, ML-BOMs aim to offer a formal way of **verifying and reporting which models are in use** and where they originate. However, we expect this category to eventually be absorbed into broader SCA functionality, as indicated by Snyk and other ASPMs adopting the functionality.

Several other vendors have also been experimenting with AI-driven posture management. **Rad Security**, **Codacy**, and **JIT**, all of whom already had established platforms, were quick to integrate AI technologies into their infrastructure and vulnerability detection workflows. These efforts typically involve building **automation around AI model discovery** and embedding **AI signals into existing posture workflows**.

Among all vendors in this space, **Noma** stands out as the most forward-looking and **innovative**. Security teams today have **very little visibility into machine learning workflows** and the tooling used daily by data engineers. Noma is built to secure the **entire data engineering workflow**, offering protection for tools like **Jupyter notebooks**, **Airtable**, and the many orchestration mechanisms data engineers use to move and transform data. This approach has the potential to evolve into a specialized version of **DSPM (Data Security Posture Management)** and **ASPM (Application Security Posture Management)**, tailored to the unique ways data engineers work.

Runtime Application Security

Application Protection

Operant

LAKERA

Enkrypt AI

NOMA

Pillar

HIDDENLAYER

splx^{AI}

PointGuard AI

From an engineering perspective, **application usage of large language models (LLMs)** is clearly emerging as the primary source of long-term risk. This has unexpectedly led to a resurgence of **RASP-like solutions**, where developers import an SDK and wrap their AI calls to enable **runtime monitoring and protection**.

Pillar Security is one clear leader in this space, primarily because it's built as a **developer-first tool**. It starts by integrating with GitHub for **discovery**, then moves to **testing** based on runtime insights, and finally offers **protection** through its SDK.

Lakera also made a strong entrance with their widely-shared **Gandalf prompt injection testing game**, and have since evolved their platform into one of the most prominent **runtime protection** solutions for AI application security.

This space is expected to become increasingly competitive as **ADR (Application Detection and Response)** providers expand into runtime AI protection. These platforms already collect much of the data needed for monitoring AI usage at runtime, and in many cases, securing AI usage is not much different from securing any other form of application behavior.

Among them, **Operant** stands out as one of the most robust runtime protection solutions. It offers both an **SDK-based** and **traditional integration** model, enabling features like **in-flight sensitive data redaction** and full coverage of runtime application behavior.

Incumbents (ADR)



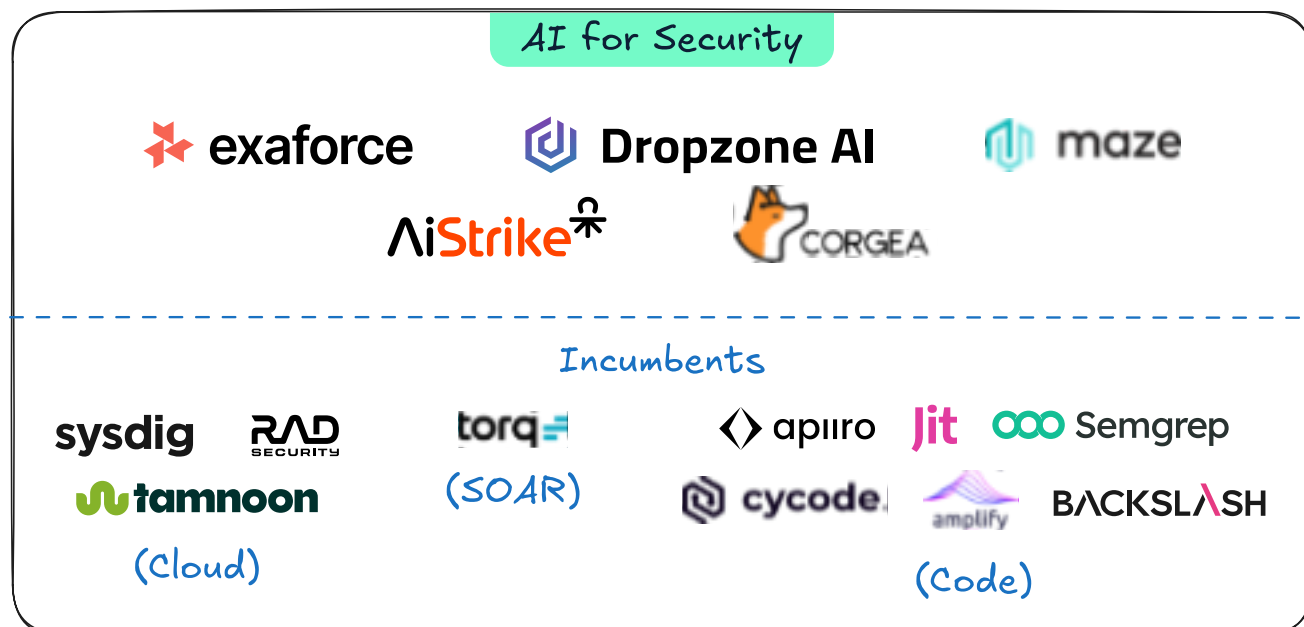
Emerging ADR providers are also able to provide similar features, without the SDK

- **Oligo** has disclosed early vulnerabilities in PyTorch and AI workflows, issues uncovered through its own monitoring tools by watching for suspicious AI package activity.
- **Raven** and Oligo, both focused on **library-level monitoring**, are best suited to securing the models themselves, particularly through deep visibility into AI-related packages and dependencies.
- **Miggo** and **Operant**, on the other hand, offer broader coverage with **multiple integration points**, enabling insight into **data flows, API calls**, and not just library usage.

As the market evolves, runtime protection will likely become a core part of enterprise AI security strategy, and the winners will be those who can balance deep model awareness with practical application-level observability.



AI for Security



Discussing AI developments in security tooling more broadly, there are two main categories worth highlighting. The first is what we'll refer to as **AI-native capabilities**, tools that have been built from the ground up with AI tightly integrated into the product itself.

The best example of this is **Corgea** (alongside DryRun and ZeroPath), who, unlike many other SAST providers that are experimenting with AI features like autofix, have **rebuilt scanning engines** to use large language models (LLMs) for the **entire detection process**. While there may be some surface-level feature overlap with providers like **Semgrep**, which also supports autofixing, the **core scanning architecture** in Corgea represents a **fundamental shift** and an evolution of the category.

Another example of this dynamic can be seen by comparing **Exaforce** to platforms like **Sysdig** or **Torq**, both of which use AI to enhance various **security operations workflows**, such as alert summarization and automated analysis. These capabilities, while basic, are already useful for practitioners embedded in those ecosystems. Exaforce, however, offers a **deeper, more comprehensive AI experience**, one that extends into **detection engineering, data querying, and incident response orchestration**.

All of the providers above are innovating with AI, but demonstrate the difference between AI-Native and AI Innovative approaches.

Vulnerability Management



A final category worth mentioning is the way **vulnerability management** is evolving with AI. **Opus Security (Acquired by Orca)** has long supported remediation workflows, but now augments almost all of them with AI to assist with prioritization, triage, and automation. On the other hand, **MazeHQ** is taking a radically different approach, where **AI handles every step** of the vulnerability investigation and remediation process, from discovery to prioritization to response.

The most exciting thing to watch is whether we'll see a **clear divide emerge between tools built before and after AI**, much like we saw with the shift to cloud-native tooling.

While the technical architecture may not change as dramatically this time, it's increasingly clear that the **user experience and data interaction model** has undergone a significant transformation.

Many vulnerability teams are also experimenting with AI for advanced vulnerability analysis. Such as Phoenix Security tying CVEs to CWEs for exploit prediction, or various academic papers published on the subject.

AI Leaders and Innovators

<i>Latio</i> AI SECURITY LEADER 2025			
			
			
			

<i>Latio</i> AI SECURITY INNOVATOR 2025			
			
			
			
			
			
			
			

TOOL BUYING FLOW CHART

Making a Buying Decision

The core question underlying any tool-buying decision in AI security is simple: how much risk does AI usage introduce to your company?

This is a particularly difficult risk to assess because AI-related risk tends to shift rapidly, from negligible to substantial, with just a few small but meaningful architectural changes. If users are primarily interacting with standard chatbots that have secure settings enabled, the risk is relatively low. These sessions are typically limited to the local context provided to the model, which simply responds to input without accessing external systems or data.

However, **risk begins to escalate quickly** once AI agents are allowed to **look up additional data**, and even more so when they're granted the ability to **take actions on behalf of users**.

Once these capabilities are introduced, the **risk profile changes dramatically**. A simple authentication misstep can allow users to access each other's data, exposing sensitive information that should remain isolated. Authentication and access control for AI agents remains an evolving challenge for many organizations, including even the most advanced model providers.

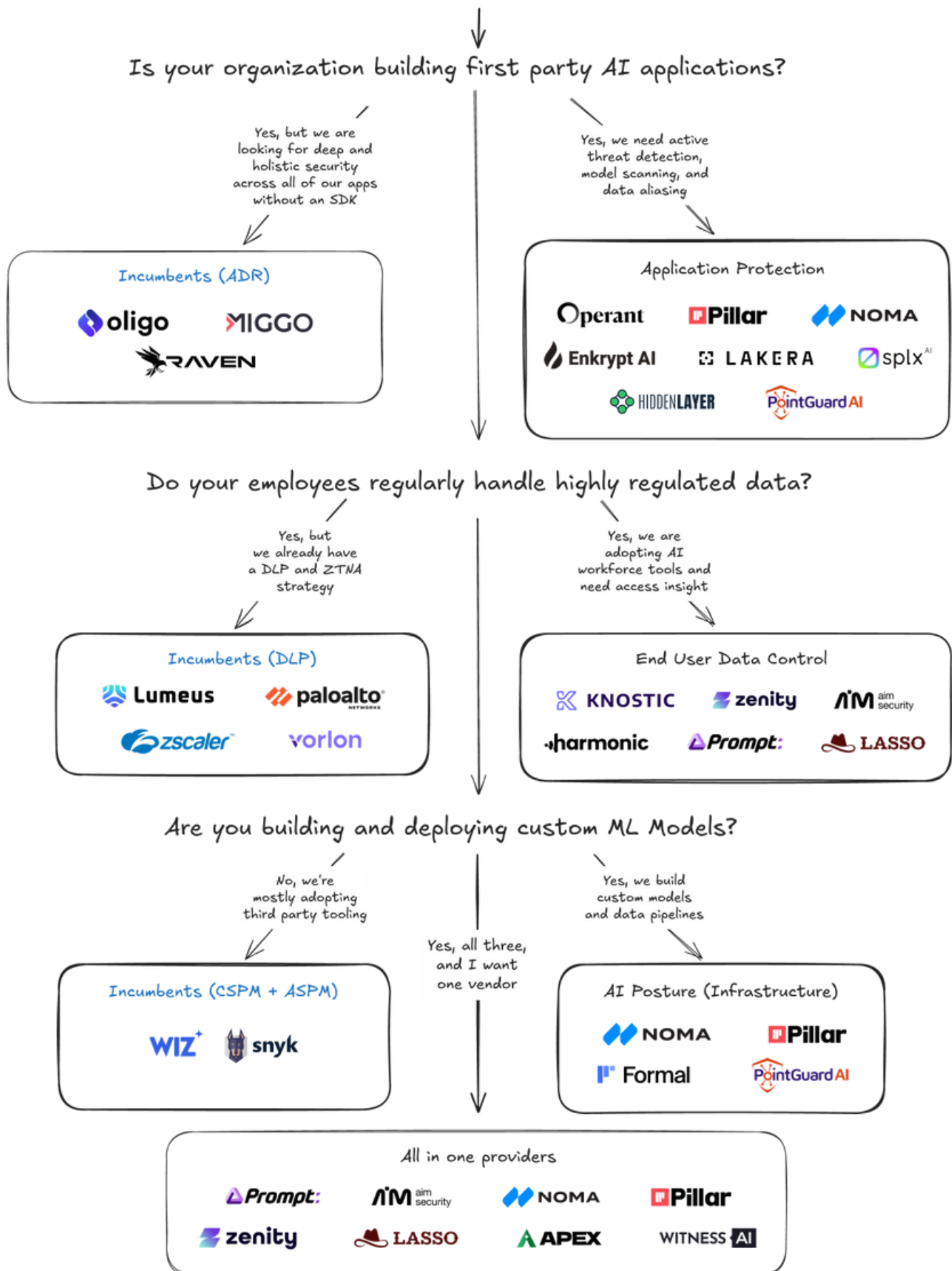
Another factor to consider is the **pace of innovation** from major platforms like OpenAI and Anthropic. These tools are rapidly adding features, including **basic security guardrails**, which, while not as robust as those from dedicated providers, indicate a growing focus on security at the model level.

So, when should you buy a **dedicated AI security solution**? The answer: **if you're moving quickly on AI adoption and handling sensitive data**, dedicated tools can offer the visibility and control your security team needs to stay ahead of emerging risks.

On the other hand, **most incumbent vendors are adding AI-related capabilities quickly**, especially for more general-purpose use cases. If your organization is not on the bleeding edge of AI adoption, it's entirely reasonable to **wait for your existing providers** to deliver the functionality you're looking for, especially if you already have a solid vendor relationship in place.

In the chart below, we've outlined a simple **decision flow** to help determine when to choose a **dedicated provider** versus an **incumbent vendor** adapting their platform for AI use.

Decision Flow Chart



Reflections on the Market

Here are some general takeaways:

AI security is not really a standalone category, but a gluing together of use cases that intersect with nearly every aspect of modern enterprise security. From data loss prevention to application protection to infrastructure posture, new risks are emerging just as quickly as new tools are being introduced to tackle them.

What makes AI security unique is how rapidly the attack surface can change. A single design choice, like giving an agent access to internal data or enabling it to take action, can turn a low-risk scenario into a high-stakes one. This variability means that security leaders need to stay closely involved in AI adoption conversations, especially as teams move from experimentation to production.

The good news is that many of the foundational techniques for mitigating AI-related risk already exist. Existing tools can often be extended to cover early use cases, while dedicated solutions offer specialized depth where needed. Choosing between the two depends not only on your current architecture, but also on how aggressively your organization is deploying AI.

Ultimately, the organizations that thrive in this new landscape will be the ones that treat AI security as an extension of their broader security strategy, built on visibility, informed by context, and ready to evolve as the technology does.

We believe the DLP use case will be the shortest-lived from a market perspective. This is why almost every vendor that started in this category has since pivoted into runtime application protection by extending their endpoint detections via API. A lot of fear, uncertainty, and doubt drove the initial adoption of endpoint-oriented tools as CISOs did not trust tooling like ChatGPT. As the market has consolidated into a few chat services, including workspace tools providing them natively, the trust is not as large an issue as it once was. Aliasing PII from system chats never made a lot of security sense.

The risk related to AI applications at runtime increases daily, as systems become more complex and users demand that their tools include sensitive data lookups. There is an interesting wrinkle here in how the industry has had an unexpected return of Runtime Application Security Protection (RASP), as most tools require an SDK to get the insights necessary to check for malicious AI usage. This puts existing ADR vendors in a great position to compete.

The rise of AI-assisted coding has also increased demand for two functionalities: detecting malicious dependencies that are doing typosquatting, and getting visibility into developer IDEs. This will lead to a long-term increase in security teams expecting their SCA tools to have endpoint integrations and detect malicious packages.

Finally, it seems likely based on early “AI-Native” prototypes that we will see a shift of tooling similar to pre-cloud and post-cloud happen. Applications built from the ground up for LLM usage function completely differently from traditional ones, similar to how containerized applications are usually built and deployed much differently than those built before the cloud.

There is an open question of how much providers such as OpenAI and Anthropic will build themselves, creating a similar dynamic to the early cloud days, where security tools were hesitant on how much to build based on what AWS would build in themselves. How much security teams need to invest in these early tools are largely based on their organizations risk profile and how quickly AI adoption is occurring.

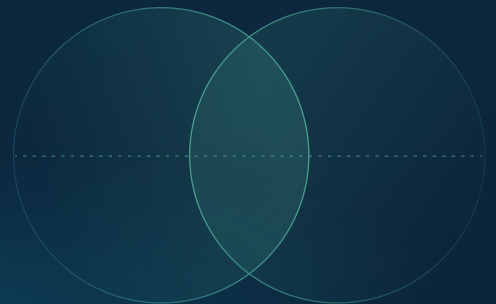
Overall, it’s apparent that AI is here to stay, and the security concerns are just getting started.

Latio

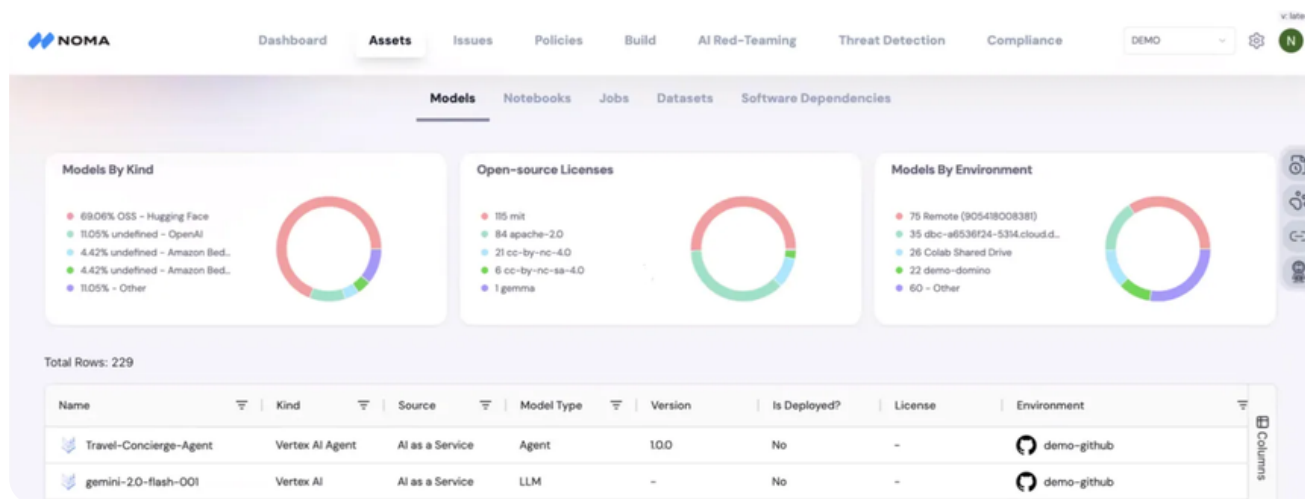
Vendor Breakdowns

Q2 • 2025

Prepared by
James Berthoty



While most of the market has been focused on browser plug-ins or SDKs, Noma is the only tool on the market that has built a holistic platform for managing the security of AI workflows across code, infrastructure, and runtime.



Noma provides a platform to secure AI agents, whether deployed to your production infrastructure or defined as code. It starts by integrating and discovering all of your AI models, from Microsoft Co-Pilot to production applications, and then allows policy creation, AI red-teaming, and threat detection across all of those models.

Most security teams are prioritizing getting visibility into their ML workflows and AI usage at a holistic level, and this is exactly where Noma started. From a feature perspective, they also support holistic application security use cases such as model scanning, dynamic testing, and runtime protection.

Why Noma is a Leader



Agentless Integrations

Integrate with your code, cloud, and SaaS providers to discover models and posture issues.



Remediate Risk

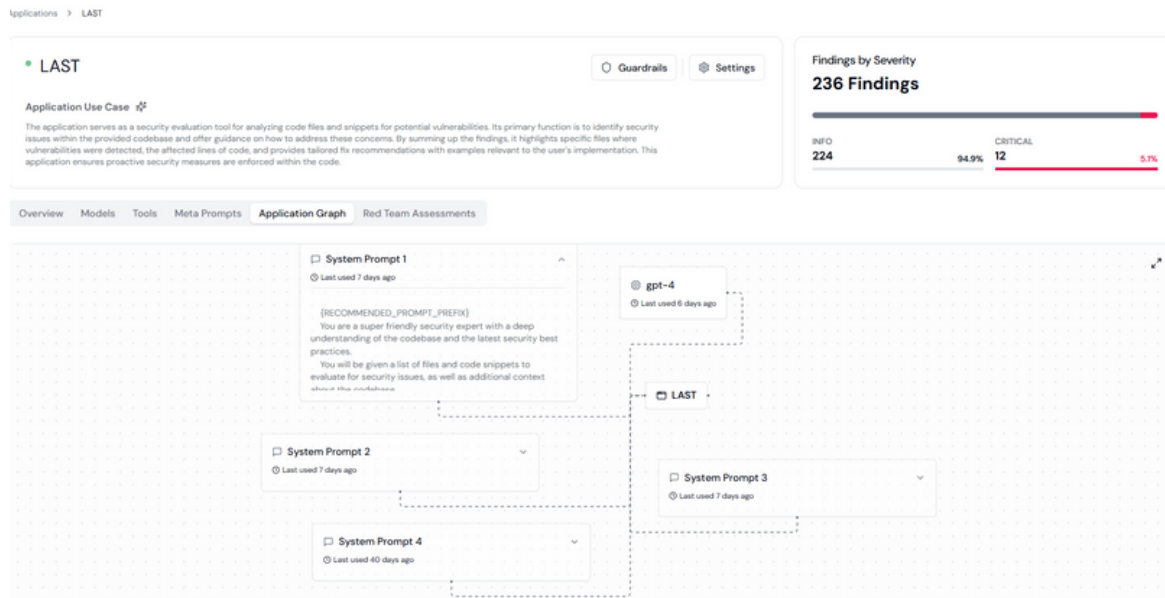
Discover and remediate configuration and model security issues.



Test & Protect

Test your deployed models and setup runtime protection guardrails.

Every security team needs ways to confidently assess and mitigate security risks across the entire AI lifecycle. Pillar Security is quickly becoming a leader in end-to-end protection for AI applications due to offering a full suite of application protection features, from mapping agentic interactions to protecting them at runtime. They offer security teams an easy way to: Discover active risks, implement a modern AI-SPM program, and AI monitoring capabilities to maintain compliance. From discovering agentic architectures to offering runtime protection, Pillar is an end-to-end application protection platform.



A real-time view from the Pillar platform of the Latio Application Security Tester (LAST)

Pillar's platform is designed for organizations that leverage agentic AI in their applications. Its uniqueness comes from the ability to build an accurate map of agentic connections and then protect those agents from attacks at runtime. Their red-teaming capabilities are similarly robust, offering a full suite of testing that is environment-specific, customizable, and powered by their threat intelligence feed. With Pillar, teams can confidently and easily run tests against different underlying model providers and leverage AI to write additional custom test cases.

Why Pillar is a Leader



Discover

Accurately discover and assess your company's usage of AI tooling, MCP servers, coding agents, models, data sets, and associated risks to proactively identify unintended AI threats.



Test

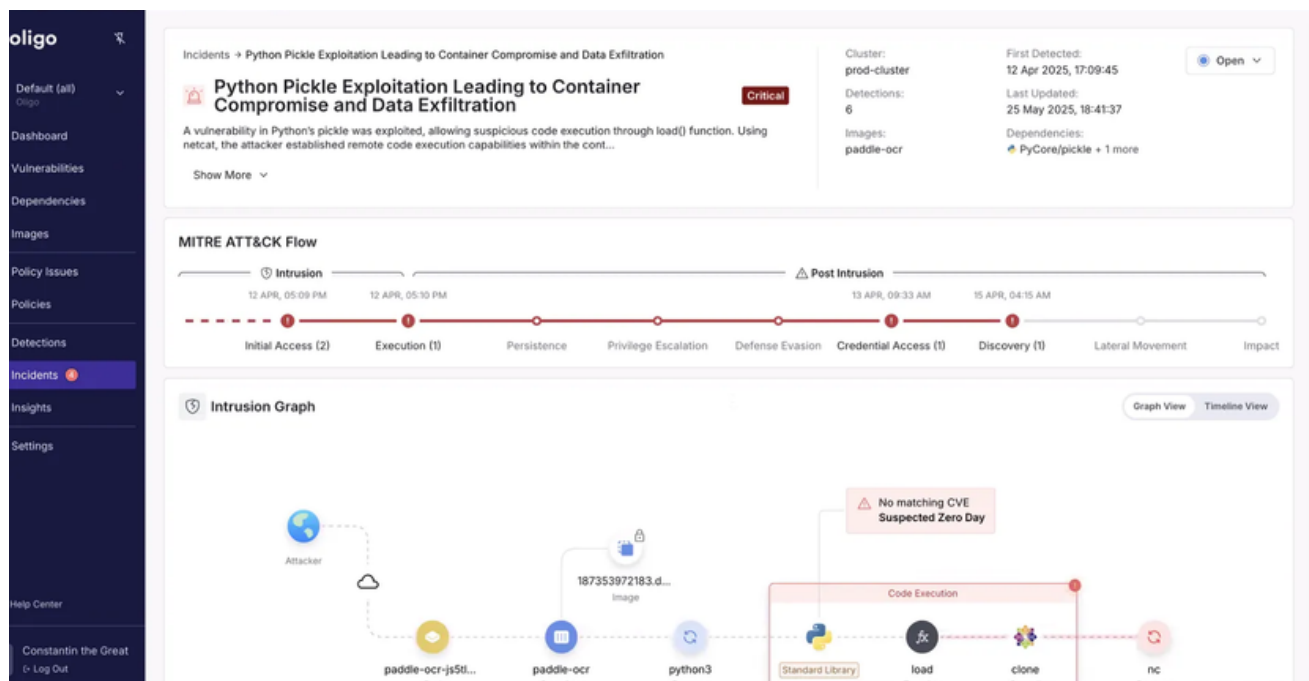
Run continuous and comprehensive tests on all deployed AI applications to detect for known and unknown active vulnerabilities.



Protect

Implement 24/7 monitoring to block active attacks at runtime, prevent data leaks, and protect sensitive data from misuse.

Oligo Security perfectly exemplifies why Application Detection Response (ADR) is essential to securing the future of all kinds of application development, with AI at the forefront. Oligo's extensive discovery of major AI vulnerabilities such as ShadowRay and a vulnerability in Llama exemplifies why robust function-level protection extends to AI frameworks as much as any other application.



Most AI work happens not only through API interaction with cloud-hosted models, but also through a lot of hosted technologies and frameworks such as Ray, or using Python's Pickle module for object serialization. Oligo is one of the only providers with deep visibility into these technologies that can watch for malicious zero-day exploits.

For example, LLM's can be exploited to accomplish remote code execution, compromising entire systems. Through Oligo's package and function-level baselining, they can immediately spot and block these sorts of attacks from occurring, in addition to helping you prioritize your vulnerabilities.

The Two Key Reasons Oligo is an Innovator:

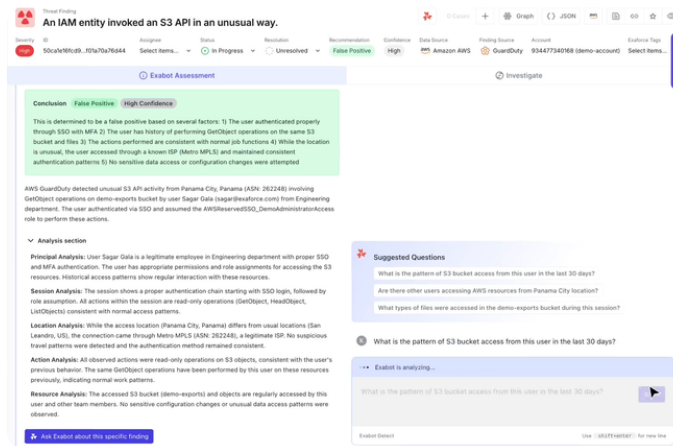
1.

It accomplishes its function-level visibility with eBPF rather than requiring SDK installations, which require time-consuming code-level integrations for developers

2.

It provides holistic application protection rather than only AI protection

Teams today are struggling to optimize their use of AI for SOC efficiency without creating data problems along the way. To solve this, a solution needs to be able to standardize and gather both posture and runtime data to enable full AI optimization. Exaforce is tackling this issue head-on. This solution is taking an “AI Native” approach to security operations by building the first end-to-end SOC platform for agentic AI. Unlike most existing players who are solving only small pieces of SOC optimization, Exaforce optimizes just about every area of the SOC with a massive data platform built for agentic optimizations, and offers MDR services alongside the platform.



The Exaforce SOC platform tackles several major challenges in SIEM architectures to build a data lake that maximizes the benefits of agentic AI’s capabilities. It helps teams by:

1. Tying identities back to actions across log sources to understand user sessions end-to-end without the heavy querying overhead.
2. Standardizing log sources and asset data to combine posture and runtime information to expedite investigations by centralizing evidence in a central location.
3. Tying together alerts from different applications to create a true XDR experience to create a comprehensive view into an attacker's actions within an environment.
4. Using agentic AI to not only summarize issues, but also take actions on behalf of users.

In my opinion, Exaforce is a glimpse into the evolution of SIEM technologies we’ve been waiting for, and in the meantime, is a massive optimizer for existing SOC spend.

Why Exaforce is a Leader



AI Analyst

Instantly analyze potential attacker activity, and validate findings from true to false positives.



AI Detection

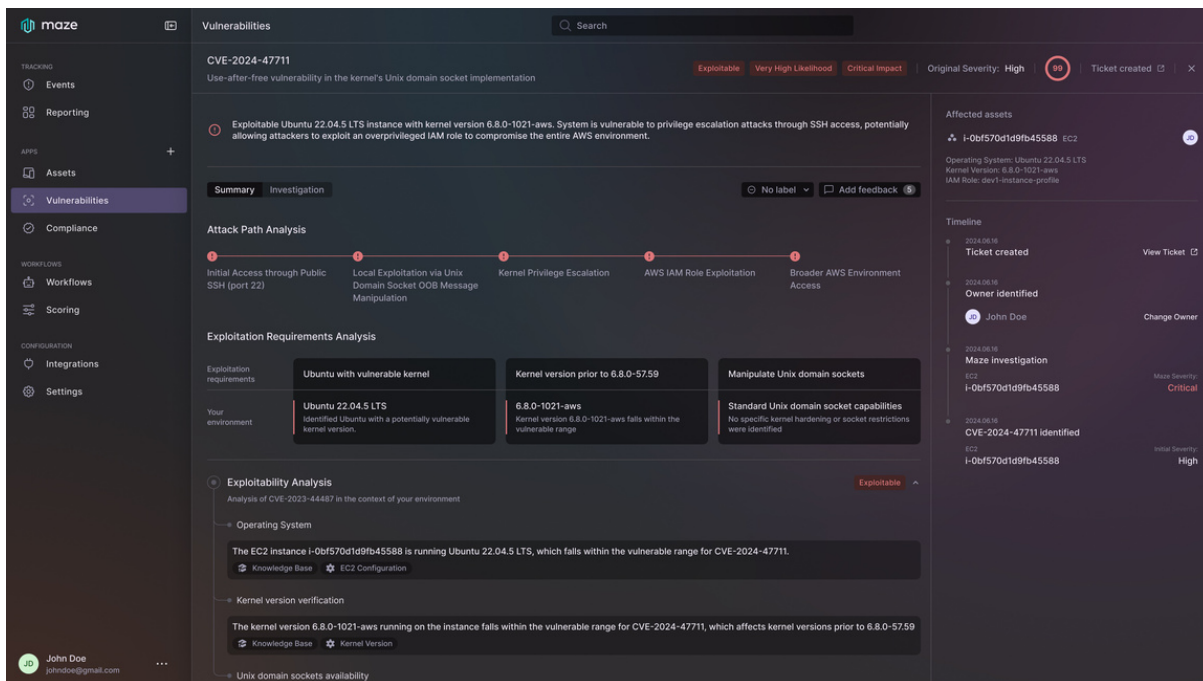
Customize rules based on detection trends across sources to reduce mean time of rule enforcement from weeks to minutes.



AI Response

Automate environment specific response workflows every time - and ensure the right stakeholders are up to date with key insights.

Maze is a major disruptor in the vulnerability management and CTEM ecosystems by taking an AI Native approach to the problem. Maze takes a deceptively simple approach to vulnerability management: letting AI investigate vulnerabilities the same way that humans do. In practice, this means taking a vulnerability finding, figuring out what exploit scenarios look like, and then investigating it in the exact context of your own infrastructure and compensating controls.



Maze starts by researching the vulnerability and understanding the full suite of available exploits, configurations, and mitigating factors that make it exploitable or not. From that deep analysis, the AI agents then investigate the configuration of your existing workloads and infrastructure to determine if the exact requirements necessary for an exploit are present. If the vulnerability can't be exploited, an exception is created, and the reasoning is logged for compliance. For exploitable vulnerabilities, Maze then does further investigations to create a business-relevant scenario about the potential impact, allowing the truly critical risks to be fixed first.

On paper, other tools surface similar data to create attack paths and vulnerability indicators; however, because Maze starts with a deep analysis of the vulnerability and works from that to your infrastructure, it creates a much more useful outcome where you can know with near certainty whether or not something is exploitable. Maze is the only platform that can truly handle how unique every vulnerability really is.

The Maze Agentic Workflow



Investigate the specifics of each vulnerability and the exploit



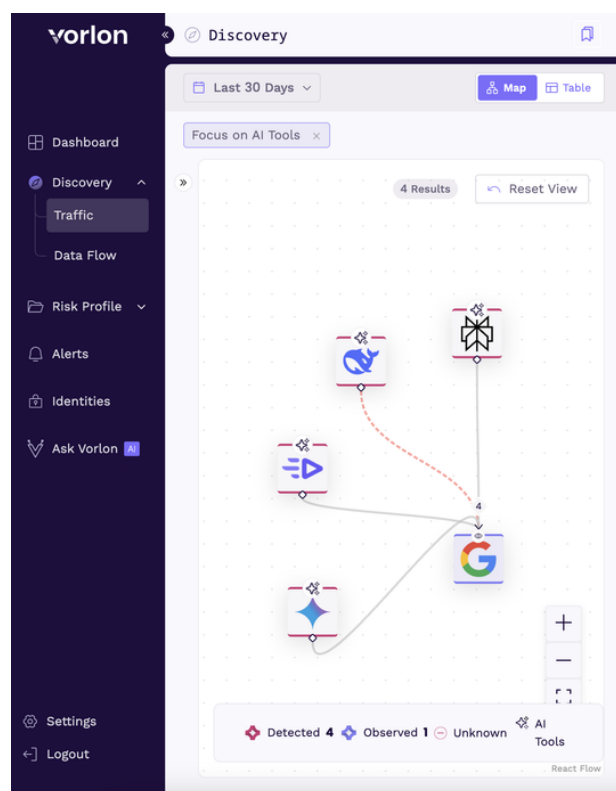
Generate an attack path from the internet to the exploitable vuln



Quickly determine if the vulnerability is a false positive or needs to be patched

Vorlon is an excellent example of how modern SaaS security platforms with runtime capabilities are expanding into runtime protection of AI applications. Runtime protection for SaaS and AI tools is inherently more challenging because security teams have less visibility and control over third-party apps. However, Vorlon's algorithmic modeling of SaaS environments helps overcome many of these visibility challenges.

As the industry consolidates into a few dedicated AI providers, it's clear that most AI tools protection will involve managing them as any other SaaS application - both getting posture-related configuration data, and runtime detections of potential compromise.



Vorlon's runtime-focused detection has two major advantages:

1. Tracking real-time communication between SaaS providers and connected services, including sensitive data flows
2. Visibility into users and NHI connections across the SaaS ecosystem, including AI tools

By integrating and monitoring traffic between these systems, Vorlon offers a comprehensive view of your SaaS ecosystem, and AI is a natural extension of those capabilities. For AI specifically, Vorlon discovers shadow AI usage, shows where sensitive data is being shared with models, observes how AI applications connect back into your overall data stack, identifies unusual data sharing patterns, and uses this intelligence as context to prioritize remediations.

Why Vorlon is an Innovator



Secure AI

See and secure how AI tools and agentic systems are accessing sensitive data



Secured MCP

Vorlon's MCP (Model Context Protocol) server enables secure communication with AI models and uses AI to detect and respond faster

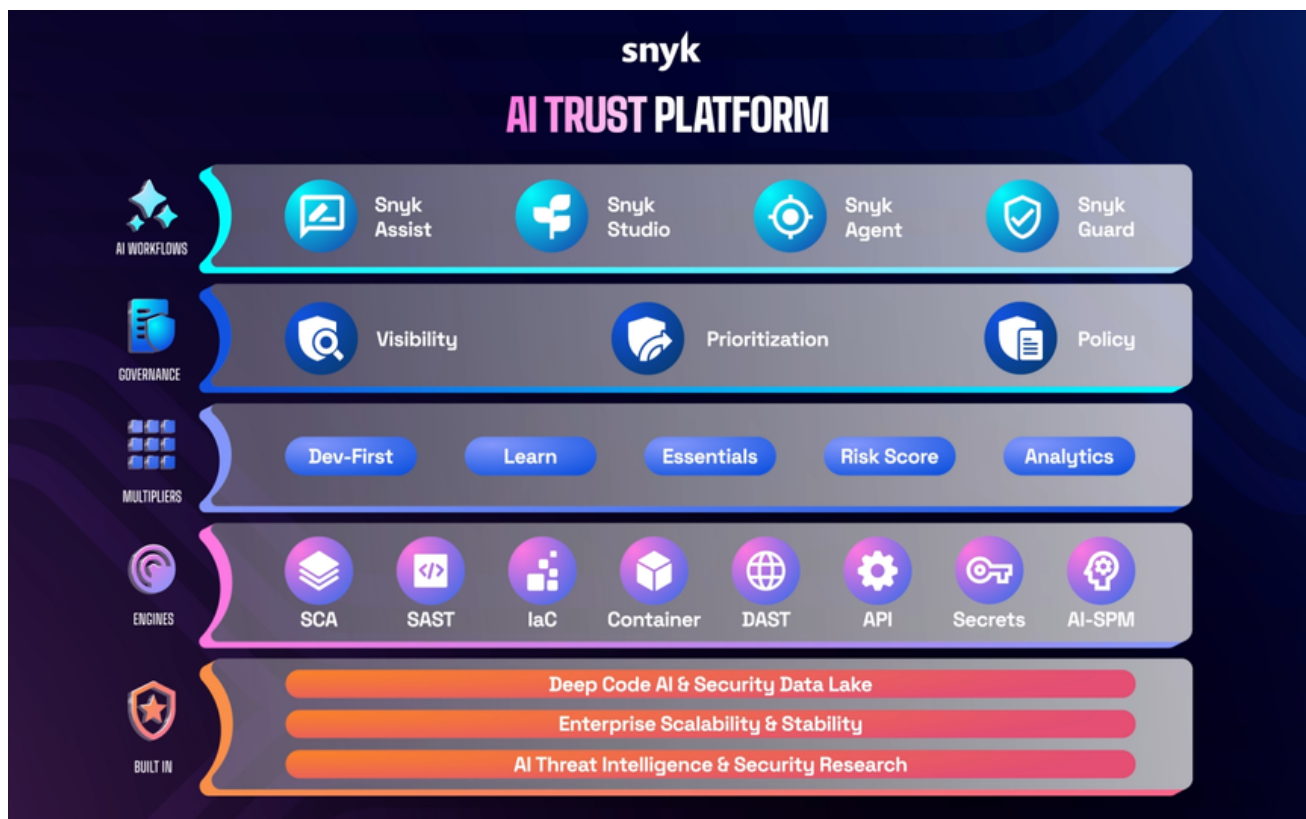


Faster Fixes

AI-powered detection and remediation guidance for fixing issues across your SaaS ecosystem, including AI tools



Snyk has recently introduced a slough of AI features aimed at both securing AI usage and accelerating application security effectiveness with AI powered tooling. These features cover everything from governance capabilities to better automated fixes.



These features cover several areas, but most important to developers are: AI-BOM and Model Risk Registry capabilities for assessing model risk, extending DeepCode AI autofixes across developer tooling from IDEs to pull requests, and the deployment of MCP servers for security scanning in AI IDE tooling. Additionally, Snyk Guard indicates Snyk's stepping into the runtime AI security arena with real-time guardrails for protection. This set of tooling is aimed at enabling enterprise application teams to adopt AI developer tooling with minimal risk.

Why Snyk is an Innovator



Increase Efficiency

Allow developers to safely use AI tooling to increase productivity, and spend less time on manual fixes.



Safe AI Adoption

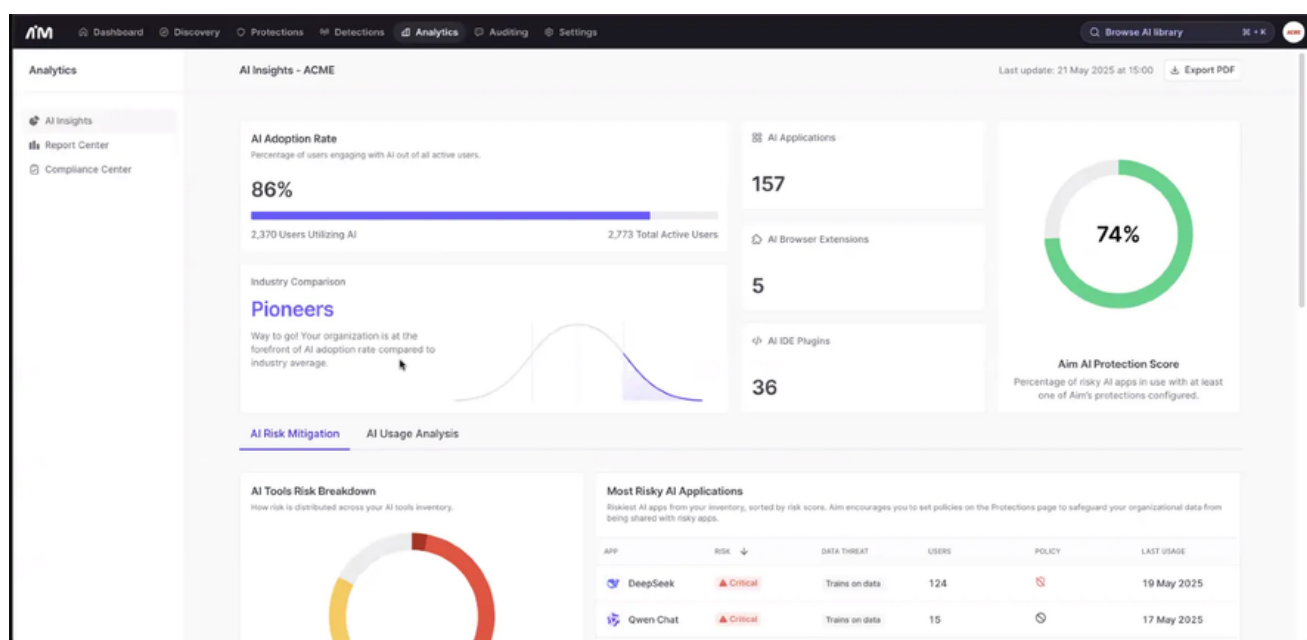
Gain visibility into your AI models, and the components and datasets that make them up.



Stay Ahead in AI

Leverage the latest in AI security research and developer security training.

AIM security provides a holistic platform empowering enterprises to accelerate AI adoption across their workforce and home-built applications. They cover a comprehensive set of use cases - from discovery via client-side installations, to API based application protection, with unified detections, reporting and data analysis across platforms.



AIM is well constructed to help enterprises gain visibility and protection as they adopt AI across their business. For large enterprises, controlling the flow of sensitive data is paramount, and GenAI applications are a sprawling opportunity to lose visibility into where your proprietary information is going. Teams are rapidly adopting a ton of client-side tooling - from tools like Cursor, to Microsoft Co-Pilot - new tools are gaining traction at a non-stop pace.

AIM's platform is unique in how unified the experience is, for example, building dataset visibility and tracking access for both Microsoft Co-Pilot and Azure Cloud. That visibility follows through to offering solutions, such as detecting that a model has vulnerabilities and access PII, and then suggesting data anonymization and guarding solutions to enforce protections.

Why Aim Security is a Leader



Discover

Utilize diverse integration types to gain full visibility into AI usage across applications, IDEs, and browsers.



Detect

Customize rules based on detection trends across sources to reduce mean time of rule enforcement from weeks to minutes.



Protect

Use pre-built or define custom detections in natural language to build and test policies for alerting, bias, and data masking



Since the last report on [AI Autofixing](#), my excitement for the future of SAST powered by AI, as demonstrated in [Corgia](#), only grows. The future of AI testing will be powered by an “AI Native” approach; in other words, using LLMs to do the contextual analysis that traditional SAST scanning even tries to accomplish.

File Path	Line Number	Urgency	Language	SLA Status	Auto-Fix	Confidence	Policy	New Finding
insecure-java/src/main/java/com/example/catapp/controllers/CatPictureController.java	39	High	Java	Due	Fix Available	Unknown		
insecure-java/src/main/java/com/example/catapp/controllers/CommentController.java	24	High	Java	Due	On Hold	Unknown		
insecure-java/src/main/java/com/example/catapp/controllers/CatPictureController.java	37	High	Java	Due	Fix Available	Unknown		
insecure-java/src/main/java/com/example/catapp/controllers/CatPictureController.java	36	High	Java	Due	Fix Available	Unknown		
insecure-api/database.py	98	High	Python	Due	On Hold	Unknown		

Issue Severity

Severity	Count
Total	39
Critical	10
High	24
Medium	5
Low	0

Fix Status

Status	Count
Fixes Available	27
False Positive	0
Processing	0
On Hold	12
Unsupported	0

The majority of OWASP Top Ten vulnerabilities at this point are some form of access control violation, contextual to your application. For example, in my testing application, I allow anyone to delete posts instead of only registered users to delete their own posts. These are the kinds of vulnerabilities that cause massive business harm, and yet go completely undetected by traditional SAST tooling.

Rules-based approaches to code analysis are failing security teams. On the one hand, you have massive numbers of false positives. On the other hand, you have a completely missing category of findings related to being contextual to your app. Corgia deftly tackles both issues.

Another example I love is a ransomware script without any SAST findings - I do this to test for malware detection that's not signature-based. Only tools like Corgia call this out, first notifying of the script's presence, but then realizing the repo is a testing repo and lowering the urgency.

Why Corgia is a Leader



AI Detection

Use AI to detect the critical application weaknesses that go undetected by traditional SAST



Contextual Prioritization

Use AI to accurately prioritize issues in the context of your overall application rather than relying on static scoring.

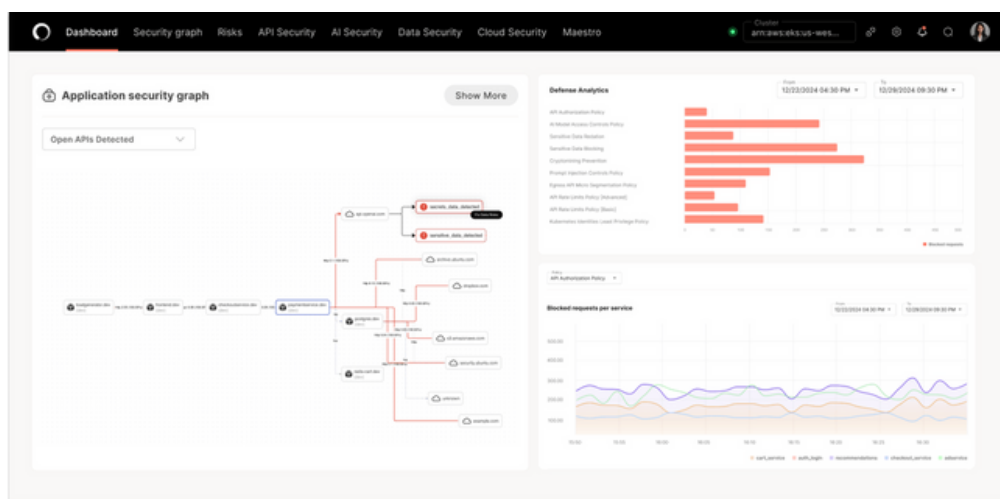


Reduce False Positives

Use AI's analysis to reduce false positives by validating findings in the context of your overall application's architecture.

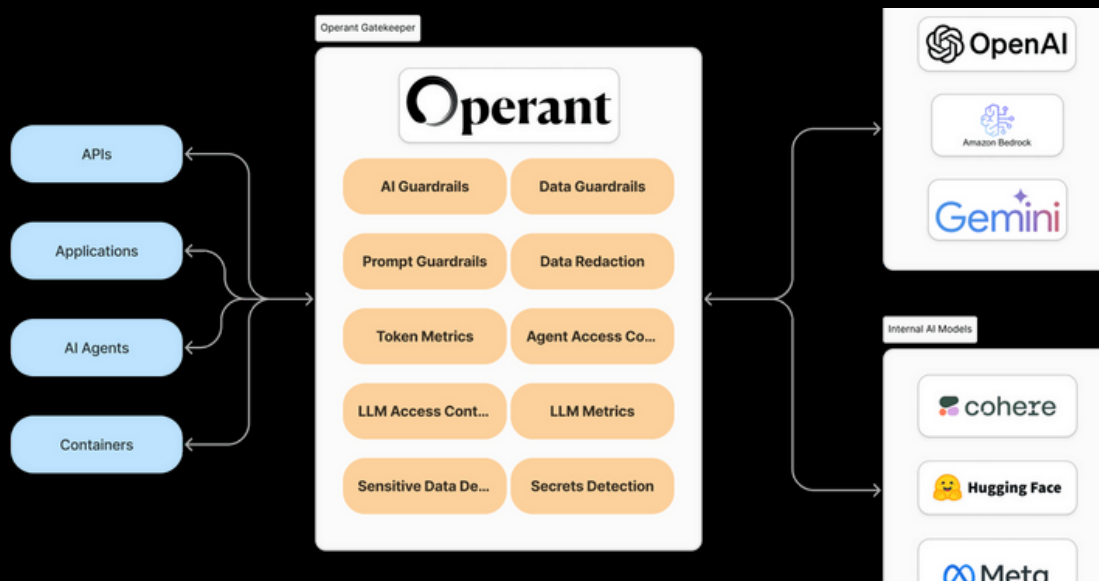


Operant stands out for being a solution in this category I'm confident actually works beyond the marketing pages. Operant has a strong background in CADR runtime protection - from API's to workloads - and has leveraged that experience to quickly become a leader in runtime security for AI applications. Operant can do in flight data redaction, and offers deep visibility and protection.



Due to their ability to function as a full runtime security solution, for AI specifically, Operant can protect just about any configuration of LLM you may be using for your applications. Whether you're using RAG or MCP, Operant is able to give visibility and guardrails every step of the way as your application reaches out to vector databases and LLM as a service providers.

Operant's approach to runtime AI application security is truly holistic by defending APIs, applications, agents, and containers all at the same time. The team's rapid expansion of CADR capabilities into AI is remarkable.



Thank you!

Thank you for reading this report! We're excited to continue delivering high quality **actually useful** product assessments that go deeper than any other reports in the industry. Latio delivers a report quarterly, including Market Reports, Category Guides, and Testing Reports.

Your support is what makes it all possible!

Follow our work at <https://pulse.latio.tech> or browse the full catalogue of vendors at <https://latio.com>

 james@latio.com

 latio.com

 Raleigh, NC



Latio

