

WIZ[★]

AI-SPM Buyer's Guide



Table of Contents

Introduction	3
Why AI security matters (and why it's challenging)	4
Common security risks in AI environments	5
Classic threats with added complexity	5
Unique AI-specific threats	6
What is AI-SPM?	6
Benefits of an AI-SPM solution	7
Unified visibility	7
Improved time-to-production for AI applications	7
Effective risk reduction	7
Streamlined incident response	8
Simplified compliance management	8
Easier collaboration	8
Self-assessment checklist for AI-SPM	8
1. Does my organization know what AI services and technologies are running in our environment?	8
2. Do I know the AI risks in our environment?	9
3. Can I prioritize critical AI risks?	9
4. Can I detect misuse in our AI pipelines?	10
Essential function checklist for AI-SPM use cases	11

Introduction

In recent years, AI's rapid growth and adoption across various industries have been nothing short of transformative. With key applications like generative AI (GenAI) chatbots, predictive analytics, and diagnostics leading the way, AI's impact is only expected to expand further, with the global AI market projected to reach over \$2,5 trillion by 2032.

As business operations, decision-making processes, and customer interactions are reshaped by AI, organizations' tech stacks are also necessarily experiencing a significant shift:

- 70% of cloud environments now utilize AI services, nearing the 81% adoption rate of Kubernetes, a foundational deployment technology.
- 53% of cloud environments use OpenAI, highlighting the dominance of GenAI in current AI adoption trends.
- 42% of organizations choose to self-host AI models, thereby needing to own more of the security shared responsibility model.

These and other statistics from our "[State of AI in the Cloud 2024](#)" report underscore AI's critical role in modern enterprises; however, they also warn us against the potential risk and losses if AI is not properly secured.

The rapid proliferation of AI technologies has introduced complex challenges, and conventional security measures are no longer adequate. Security teams are now tasked with extending their SecOps practices to develop robust AI security measures. The recommendation is to navigate this complex landscape with an agile, low-overhead, and easy-to-use approach that can start securing your AI systems now.

With AI security posture management (AI-SPM), your organization can rely on a centralized platform that enables you to evaluate your AI security posture, mitigate AI risk, and enhance security seamlessly so you can confidently innovate with AI. Below, we provide an overview of how and why AI-SPM is essential for AI security and offer practical tips and checklists to effectively implement AI security for your organization.

Why AI security matters (and why it's challenging)

AI systems are dynamic and unique in nature, and security teams tasked with keeping a secure IT environment face several new challenges:

- **AI systems are complex and can lack transparency.** AI models are non-deterministic, often function as black boxes, and come with numerous parameters trained on large amounts of highly variable structured and unstructured data.
- **AI systems learn from user data and may violate privacy.** AI models learn from large data sets that often contain sensitive information, as such data is particularly useful for personalization.
- **AI systems lack standardization and require expert knowledge to understand.** No AI application is ever the exact same, and a variety of frameworks and paths to production often need to be maintained.
- **AI systems have the potential for large-scale misuse and need to abide by various and ever-changing regulatory requirements.** AI applications are open to a variety of attack paths that are both internal (e.g., harmful outputs or misconfigured pipelines) and external (e.g., deep fakes or automated hacking).

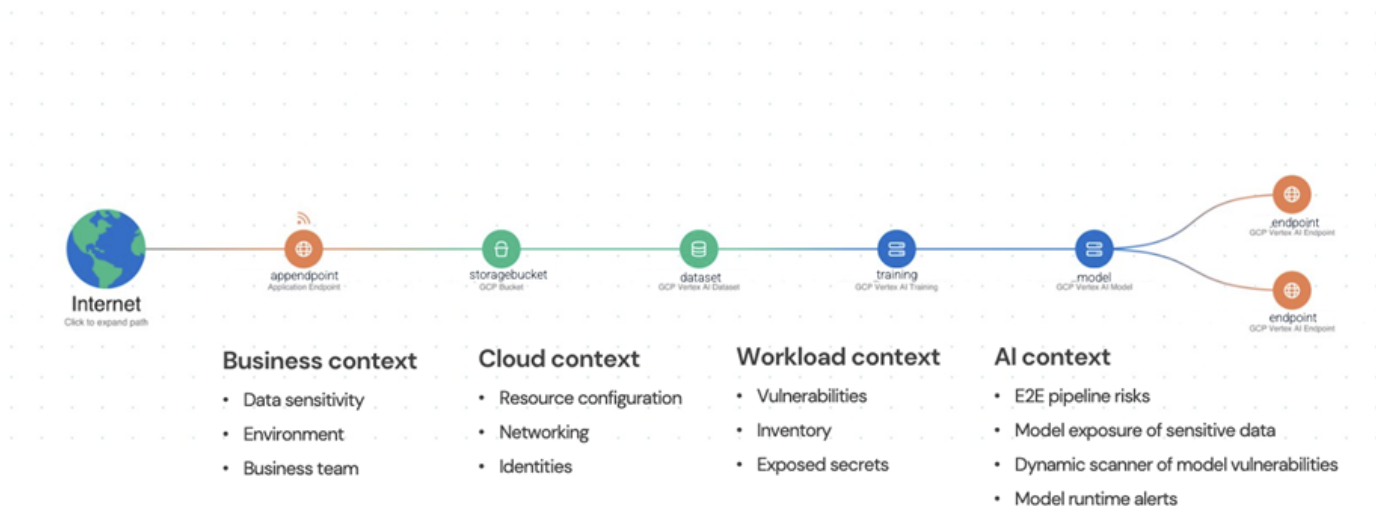
The unique nature of AI systems introduces specific vulnerabilities and evolving threats that standard SecOps practices cannot address with a one-size-fits-all approach. Addressing these challenges is fundamental to protecting not only the organization but also its customers, partners, suppliers, stakeholders, and—in some cases—broader society.

Effective AI security posture management embraces an agile mindset to safeguard sensitive data, maintain system integrity and performance, prevent financial losses, build customer trust, ensure legal and regulatory compliance (e.g., GDPR and CCPA), and preserve control over the AI system.

Common security risks in AI environments

AI environments face an array of security risks, both from traditional threats and those unique to AI. These concerns are more complex and multifaceted than typical IT threats, demanding specialized attention and mitigation strategies.

How risky is this AI pipeline?



To effectively secure AI environments, organizations must adopt a holistic approach that encompasses the business, cloud, workload, and AI contexts to ensure that all potential attack paths are identified and mitigated. This will secure the entire AI lifecycle from data collection to deployment.

Classic threats with added complexity

Traditional security threats such as malware infections, network security issues, insider threats, and unauthorized access still apply to AI systems. However, these often come with added layers of complexity in AI environments:

- **Data breaches.** AI systems are trained on large quantities of data, therefore can lead to critical data risks. These systems are vulnerable to sophisticated hacks, including model inversion attacks, membership inference attacks, model poisoning, and data extraction via prompt engineering. These and other threats, such as those presented in the [OWASP Top 10 for LLM](#), can expose sensitive data and compromise the integrity of AI models.
- **Denial of service (DoS).** AI endpoints, especially those utilizing GenAI, can be overwhelmed by large query volumes or manipulated through prompt engineering. Additionally, AI algorithms can be exploited for their computational vulnerabilities, leading to service disruptions.
- **Supply chain attacks.** AI systems frequently depend on pre-trained models, data sets, and third-party libraries. These components can be compromised via malicious code injection, backdoors in third-party services and APIs, or data sets and models that have been tampered with.

Unique AI-specific threats

AI introduces new attack surfaces and risks that extend beyond what traditional SecOps faces. These include:

- **Model theft and intellectual property risk.** AI models can be extracted, reverse-engineered, or accessed without authorization, leading to the theft of model architecture, weights, and proprietary algorithms through side-channel attacks.
- **Adversarial attacks.** Hackers can manipulate AI inputs, conduct evasion attacks, or poison training data to corrupt model outputs or behavior, undermining the reliability and safety of AI systems.
- **Shadow AI.** The rapid adoption and ease of access to third-party AI solutions can lead to unmonitored and unauthorized AI implementations within an organization, increasing the risk of security breaches.

New threats and attack techniques are continuously emerging, so it is critical for companies to stay current using resources such as [MITRE ATLAS](#).

What is AI-SPM?

[AI-SPM](#) provides visibility into AI pipelines and their security posture so organizations can effectively remove AI risks. It embeds security within your AI models, training data, and AI services, ensuring accelerated—but safe—adoption of AI technologies. AI-SPM solutions should enable you to protect your AI models by providing:

1. Visibility through AI-BOM

An AI bill of materials (AI-BOM) offers unparalleled visibility into AI assets, including managed AI services and hosted AI technologies. It helps detect and track all AI usage within an organization, removing shadow AI and allowing for real-time monitoring and identification of unexpected deployments.

Example: Detect an unexpected deployment of an Amazon Bedrock service, which could indicate unauthorized or unplanned activities.

2. Proactive risk mitigation with context

AI-SPM identifies critical attack paths to AI models and data by detecting AI risks and correlating them with deep cloud context for effective risk prioritization.

AI-SPM should be integrated within a greater cloud security platform to include context around misconfigurations, vulnerabilities, identities and permissions, network exposures, secrets, sensitive data, and malware. This allows you to protect sensitive training data, remove model vulnerabilities, detect malicious models, and more via context on a security graph

Example: Detect a misconfigured IAM setting that exposes an S3 bucket containing model weights to the public, thereby preventing unauthorized access and potential data breaches.

3. Threat detection in AI pipelines

AI-SPM actively monitors and detects suspicious activity within AI pipelines. It identifies misuse or anomalies that could indicate security breaches, allowing for swift containment and removal of attack paths. Threat-detection data is correlated back to the graph-based context to empower you to understand and reduce the blast radius.

Example: Detect an AI model that is invoked by a previously unseen country or principal, which could indicate attackers are trying to infiltrate the AI system.

AI-SPM simplifies security processes for both security personnel and AI engineers, helping upskill security teams on AI and AI engineers on security by providing accurate risk prioritization with easy-to-understand graph-based context.

Additionally, by integrating seamlessly with popular third-party tools, AI-SPM helps easily automate responses to risks with your existing tech stack and further streamline security management within existing workflows.

Benefits of an AI-SPM solution

Since AIOps is relatively new for many organizations and traditional security tools are not typically fully equipped to handle AI security, organizations face the danger of potentially leaving gaps in their AI security posture. AI security posture management addresses these unique challenges, providing a specialized and centralized platform that inherently enhances the security of AI environments.

Unified visibility

AI-SPM's single pane of glass view into the entire AI ecosystem and cloud environment, eliminating the need for juggling multiple tools and dashboards. This allows security teams to monitor AI activities comprehensively, saving time and reducing the risk of overlooking potential threats.

Improved time-to-production for AI applications

AI-SPM fosters a proactive security posture with standardized processes. These speed up the deployment of AI applications, replacing ad-hoc security checks with continuous monitoring, automated risk assessment and prioritization, and integration with existing security frameworks and compliance requirements.

Effective risk reduction

By providing continuous monitoring for AI security posture, correlated to cloud risks, security and AI teams can effectively remove attack paths to AI. This enables organizations to truly focus on the most critical risks and protect their AI models, instilling confidence in their AI innovation.

Streamlined incident response

With centralized monitoring and alerting, security teams can quickly identify and respond to AI-related security incidents. Context-rich alerts help reduce false positives and enable faster triage, ensuring that potential threats are addressed promptly and effectively.

Simplified compliance management

By providing continuous compliance assessments against relevant regulations and standards, as well as generating compliance reports as needed, AI-SPM transforms compliance management from a labor-intensive and error-prone endeavor into a reliable legal and regulatory process.

Easier collaboration

AI-SPM solutions feature a centralized platform that makes it easy to use and understand risk, even for non-security experts, enabling security democratization across the organization.

Accurate risk prioritization helps build trust and facilitates better communication and collaboration between security teams, data scientists, and other stakeholders involved in AI development and deployment.

Self-assessment checklist for AI-SPM

Organizations can evaluate their readiness and requirements for AI security posture management by answering the following questions. This checklist will help assess your current security stance and what actions are required to enhance your AI security posture.

1. Does my organization know what AI services and technologies are running in our environment?

AI-SPM's single pane of glass view into the entire AI ecosystem and cloud environment, eliminating the need for juggling multiple tools and dashboards. This allows security teams to monitor AI activities comprehensively, saving time and reducing the risk of overlooking potential threats.

Maintaining an AI-BOM is crucial for inventory management, ensuring that you have a comprehensive understanding of all AI services and technologies in use.

- ✓ Track all AI models, their versions, and associated data sources.
- ✓ Have visibility into third-party AI services and components.

Rather than relying on manual flagging, a systematic approach to the development of the AI-BOM ensures thorough and up-to-date visibility so that nothing is overlooked.

Ask yourself:

- ☐ Do we have a centralized and continuous inventory of all AI models in production?
- ☐ Can we track the lineage of our AI models, including training data sources?
- ☐ Do we maintain version control for our AI models and associated code?
- ☐ Are we aware of all third-party AI services or components integrated into our systems?
- ☐ Are we able to map our AI assets to specific business applications and processes?

2. Do I know the AI risks in our environment?

Identifying and cataloging potential AI risks is key to implementing the right strategy for managing them.

- ✓ Understand AI-specific threats.
- ✓ Identify ethical, privacy, and security risks.

Categorizing risk between data risk, model risk, and model pipeline risk ensures that you are securing your end-to-end AI lifecycle.

Ask yourself:

- ☐ Have we conducted a comprehensive AI risk assessment across our organization?
- ☐ How has the attack surface changed?
- ☐ Do we continuously monitor for new AI risks in our environment?
- ☐ Can we see what AI assets have been impacted by emerging vulnerabilities?
- ☐ Do we understand the data privacy implications of our AI systems?
- ☐ Where do we have sensitive data and are there exposure risks?
- ☐ What are the risks in the third-party AI tools we use? Are there any vulnerabilities?
- ☐ Can we detect misconfigurations in AI services?
- ☐ Can we detect malicious models deployed?
- ☐ Are we aware of the regulatory risks related to our AI implementations?

3. Can I prioritize critical AI risks?

Prioritizing AI threats allows companies to channel resources to those vulnerabilities that pose the greatest immediate threat.

- ✓ Address the most critical risks first with context.
- ✓ Balance technical and business impact.

Having a centralized dashboard for your AI risk analysis with graph-based context is a fundamental tool for accurate risk prioritization.

Ask yourself:

- ☐ Can we effectively prioritize AI risks based on business impact?
- ☐ Do I have context around AI risks to understand critical attack paths to models and data?
- ☐ Are we able to correlate AI risks to cloud context to understand attack paths from AI to our cloud environment?
- ☐ Are we involving key stakeholders from both technical and business teams in AI risk prioritization discussions?
- ☐ Can our AI risk management strategies evolve with the changing threat landscape?

4. Can I detect misuse in our AI pipelines?

Detecting misuse and anomalies in AI processes is critical, as this allows you to ensure your AI system's security and remove threats in real time.

- ✓ Detect malicious activities in real time.
- ✓ Identify the blast radius of a specific threat

The aim is to detect misuse and unfolding threats in AI pipelines, as well as understand the cloud context around them, so you can quickly limit the blast radius.

Ask yourself:

- ☐ Do we have monitoring systems in place to detect unusual patterns in AI pipelines?
- ☐ Can we identify unauthorized access or modifications to our AI models or training data?
- ☐ Do we have mechanisms to detect data poisoning attempts in our AI pipelines?
- ☐ Do we have alerts set up for abnormal resource usage by our AI systems?
- ☐ Can we detect attempts to reverse-engineer or steal our AI models?
- ☐ Do procedures exist to detect and investigate potential AI-related insider threats?
- ☐ Are we able to monitor the entire AI lifecycle for potential misuse, from data collection to model deployment?
- ☐ Can we map suspicious activity in AI pipelines to the greater cloud context to understand how an attacker could move around the environment?

Essential function checklist for AI-SPM use cases

To build a secure AI environment, ensure you can deploy complete capabilities for AI pipeline scanning, contextual understanding, analysis, and policy enforcement. This essential function checklist will help you gauge what security safeguards your organization may need to prioritize as part of your AI security posture management.

Data governance and security

- Training data security and privacy assessment
- Data lineage tracking for AI models
- Data poisoning detection
- Sensitive data exposure analysis in AI pipelines
- Access control and encryption for AI data sets
- Compliance checks for data usage in AI systems

AI development security

- AI model architecture security analysis
- Dependency scanning for AI libraries and frameworks
- Static analysis for AI-specific security issues
- Secure coding practices for AI/ML

AI DevSecOps

- AI model vulnerability assessment in CI/CD pipelines
- AI model registry scanning
- AI/ML pipeline security checks and hardening
- Infrastructure-as-code (IaC) scanning for AI deployments
- Custom security policy enforcement for AI model deployment
- Automated security testing for AI systems
- Continuous integration of security with the AI development lifecycle

Model deployment and real-time monitoring

- Continuous AI model risk monitoring
- AI system identity and access management
- AI-specific threat detection (e.g., model poisoning, evasion attacks)
- Malicious model scanning
- Contextual AI risk prioritization
- Secure model serving and API security
- Real-time model security monitoring

AI infrastructure and platform security

- Cloud-based and self-managed AI platform security assessment
- AI development environment protection
- AI model serving infrastructure hardening
- AI-enabled applications and services security evaluation
- Cross-platform AI security policy enforcement
- AI-specific access control
- Network segmentation for AI systems

AI risk management and compliance

- AI-specific risk assessment and management frameworks
- Regulatory compliance monitoring for AI systems
- AI incident response and forensics capabilities
- AI system auditing and logging
- Third-party AI vendor risk management
- Continuous AI risk monitoring and reporting

What's Next?

A proactive and agile approach to AI and GenAI security is necessary to keep up with the speed of development and adoption of these technologies.

An [AI-SPM](#) tool that supports security teams in setting up their AI posture with built-in best practices and automated processes is a big differentiator for ensuring that GenAI brings only the desired benefits to your organization.

Discover more about AI security and [how Wiz can help](#). Sign up for a demo to [see Wiz in action](#) today.

Wiz AI-SPM helps organizations accelerate AI adoption while safeguarding against AI risks by discovering AI pipelines, detecting misconfigurations, and uncovering attack paths to AI services. See how today.

[Get a Demo](#)