



Building a secure-by-design culture for cloud AI security

Wiz and Microsoft help organizations address the increasing challenges of AI security in the cloud. Together, they deliver greater collaboration between teams to mitigate risks and ensure responsible AI management.



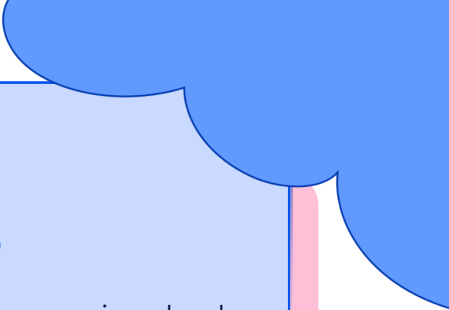
The evolution of AI adoption across organizations

Most companies across most industries in the world now leverage AI in some form, either for proof-of-concept projects or fully integrated into production environments. At the end of 2024, over 72% of companies globally have either adopted or are exploring [the adoption of AI](#). This high adoption rate is often compared to the rapid rise of technologies such as Kubernetes, highlighting the exciting potential of AI for innovation and growth.

A key reason behind the surge in AI popularity is the development of accessible AI models and services. Many organizations have implemented cloud-based services or open-source models with lower entry barriers that allow them to use AI while keeping costs down. These developments reflect the increasing maturity of the AI ecosystem, giving organizations a broader range of options that can meet their specific operational needs.

The release of mainstream AI models, such as ChatGPT, has enabled companies to quickly explore generative AI use cases. Generative AI has swept the globe, with organizations eager to use this technology to stay competitive and boost productivity — even as new security and ethical considerations emerge.

With AI models getting smarter daily, organizations must have a secure strategy to address this technology's unique security risks and challenges.



70%

of organizations are using cloud-based managed AI services such as Microsoft Azure

228%

of Azure OpenAI instances observed across all cloud environments grew between June and October 2023

32%

of organizations are still in the experimentation phase with AI cloud services

10%

of organizations are AI "power users"



Moving from AI experimentation to full-scale deployment

According to Microsoft, more companies are moving from initial innovation phases into active implementation, where they roll out generative AI applications at a large scale.

The responsibility for AI development is shifting from specialized innovation teams to core research and development (R&D) departments. This shift indicates a broader integration of AI into overall technology strategy and development. Cloud providers such as Microsoft are pivotal in accelerating this AI adoption by offering cloud services that help streamline the R&D process.

Businesses are also starting to integrate open-source AI models into their operations. They are applying the principles of scalability and economy of scale to create more cost-effective, high-impact applications.

This expansion of generative AI use within companies signals an exciting new phase in AI adoption, but organizations must understand and prepare for the risks involved.



According to Microsoft, more companies are moving from initial innovation phases into active implementation, where they roll out generative AI applications at a large scale.

Understanding emerging AI threats: is the organization at risk?

As AI becomes more deeply integrated into core business operations, new and complex security challenges emerge. AI's reliance on vast amounts of data is at the core of these challenges.

Several types of data fuel generative AI, including:



Foundational training data



**Fine-tuning data to tailor models
for specific applications**



Real-time data from external sources

These dependencies increase the risk of data leaks, sensitive information disclosure, and other critical security vulnerabilities.

As organizations rely more on AI to process and generate text based on their proprietary data, sensitive information disclosure is a significant risk.

Another significant risk is prompt injection, where attackers use natural language to manipulate AI models to disclose unauthorized information or perform restricted actions.

Top 10

The Open Web Application Security Project ([OWASP](#)) has identified the top 10 security threats specific to generative AI applications, highlighting emerging issues unique to AI.

OWASP Top 10 for 2025

- ✓ Broken access control
- ✓ Cryptographic failures
- ✓ Injection
- ✓ Insecure design
- ✓ Security misconfigurations
- ✓ Vulnerable and outdated components
- ✓ Identification and authentication failures
- ✓ Software and data integrity failures
- ✓ Security logging and monitoring failures
- ✓ Server-side request forgery

Further key risks include excessive agency, over-reliance, and expanded attack surfaces. New components such as plugins and orchestrators can execute tasks, browse the web, and extract information from potentially harmful sites. Even open-source AI models present risks, as they can be vulnerable to tampering and malicious modification.

The more a business implements AI applications across its ecosystem, the more potential AI has to autonomously execute actions that could cause harm. Leaving AI to its own devices can damage an organization without human oversight.

Due to the growing complexity of securing AI in the cloud, companies must take a more comprehensive approach to security beyond a single tool or service.

Future trends and predictions for AI security

As AI technology advances, companies must adapt their security strategy in response to its new capabilities and use cases.

Autonomous agents

One of the most exciting developments in AI is the rise of autonomous agents. These are AI systems designed to perform complex tasks independently of humans. Autonomous agents are already being integrated into business operations. They will require strict security controls to prevent them from acting beyond their intended permissions.

Sophisticated deepfakes and phishing

Malicious actors use generative AI for deceptive practices such as deepfakes, phishing, and automated social attacks. As AI-generated media becomes more sophisticated, security teams need advanced tools and processes to detect and respond to these threats.

73%

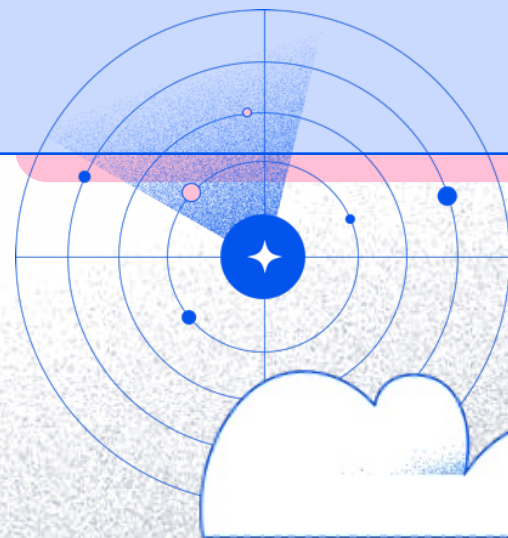
of organizations perceive significant potential risks from AI

77%

of organizations see the new risks in generative AI applications

80%

of security leaders say data leaks are their main concern



No-code platforms

Low-code and no-code platforms built with generative AI allow people without extensive programming knowledge to create applications using natural language inputs.

While the democratization of AI development is exciting, it also poses new security questions. How can organizations ensure that no-code applications are secure by design? Are there safeguards to prevent sensitive data from being accidentally exposed in these applications?

Free-tier computing

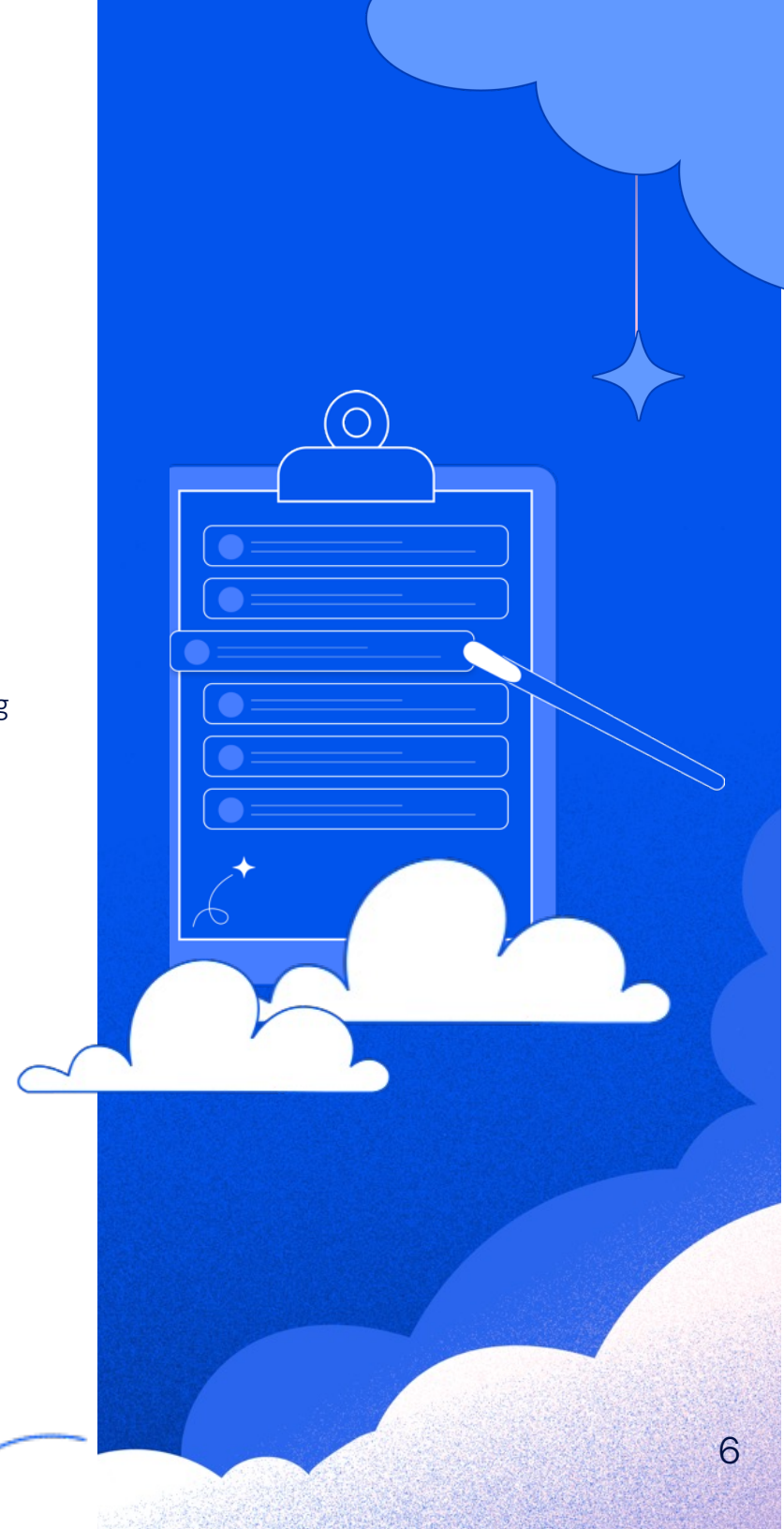
Free-tier computing services allow organizations to experiment with training their own AI models. However, they attract unwanted attention from attackers looking to exploit free computing resources for purposes such as cryptocurrency mining.

This poses a risk to users of free computing services and the cloud providers hosting them. Organizations must be prepared as attackers may use these services to attempt to move within systems to access additional resources or sensitive data.

Securing AI in the cloud

To secure cloud AI against current and future threats, organizations must start with a strong foundation of visibility and collaboration. Wiz and Microsoft advocate for greater collaboration between security teams and those developing and deploying AI systems.

AI development requires input from multiple departments, so security should be considered a “team sport” that involves both R&D and cybersecurity experts. This approach enables organizations to confidently innovate while addressing the many risks of using AI technologies.



Security teams must work alongside developers to establish security checkpoints throughout the AI lifecycle. This collaboration ensures that models are evaluated for vulnerabilities from early development through production.

Key questions for security leaders

Here are some questions to guide security teams toward safeguarding their AI systems:

Are there security controls in place to protect sensitive data?

Data is the fuel for generative AI, so protecting it is essential. Security teams should ensure data is encrypted, controlled, and managed with strict privacy protocols.

Is there visibility into the entire AI pipeline?

Security teams must understand every part of the organization's AI pipeline, from initial training data to final outputs. This visibility helps identify potential vulnerabilities, including those in third-party and open-source models.

Can the organization detect and respond to threats in real time?

Given AI's potential for rapid, autonomous actions, security teams should monitor AI applications closely and have automated response mechanisms in place.

Are AI models audited regularly?

AI models are continuously evolving, which affects their accuracy and security. Regular model validation helps ensure that models maintain alignment with organizational security policies and operate within safe boundaries.

Do AI applications respect user permissions and sensitive data labels?

Data overexposure is an ongoing risk in generative AI. Security teams must ensure AI applications are designed to respect data permissions and prevent unauthorized information sharing.

Security efforts should focus on known vulnerabilities and emerging threats. With a collaborative, proactive team approach, organizations minimize data visibility gaps in their processes and confidently use AI.

Secure-by-design: Creating greater AI security with Wiz and Microsoft

Wiz and Microsoft help organizations achieve significant security improvements in deploying and managing AI solutions across the cloud ecosystem.

 **Bridge the gap between security and development**
Wiz and Microsoft advocate for a “team sport” approach to AI security, with close collaboration between security and development teams throughout the AI lifecycle. This approach ensures that security is factored into every AI development and deployment stage, leading to more secure and robust environments.

 **Increase visibility into AI assets and activities**
With Wiz and Microsoft, organizations have complete visibility into where AI is used, including identifying AI models, services, and the data used for training. Get visibility into AI workloads and potential security issues.

 **Identify shadow AI**
Wiz’s expertise in cloud security and Microsoft’s cloud AI services allow organizations to pinpoint instances of AI in environments without proper authorization or oversight.

 **Map data flows**
By using Wiz, organizations clearly understand how data is used in AI pipelines, including training, fine-tuning, and real-time data usage. This understanding enables better data governance, ensures that sensitive information is handled appropriately, and mitigates risks of data leaks.

 **Monitor AI usage**
Wiz provides insights into which AI models are deployed in organizations, where they are located, and how they are used. These insights are essential for tracking model performance, identifying potential vulnerabilities, and detecting any unauthorized or malicious use of AI models.

 **Protect sensitive data**
Microsoft provides a secure cloud foundation with services such as Azure Security Center and Microsoft Sentinel to protect sensitive data in the AI ecosystem.



Mitigate emerging threats

Wiz helps organizations secure everything they build and run in the cloud. With Microsoft, organizations get a multi-layered defense against current and emerging AI threats.



Reduce the risk of excessive agency

Wiz and Microsoft enable human oversight and clearly defined boundaries for AI. With Wiz and Microsoft, organizations enjoy all the benefits of AI while ensuring human control over critical decisions and actions.



Promote secure-by-design principles

Wiz and Microsoft focus on AI specific security controls that enable organizations to build secure AI systems from the ground up. Doing so minimizes vulnerabilities and reduces the risk of security incidents.



Take a more proactive security posture

Wiz and Microsoft enable organizations to establish a proactive security posture for building and running AI models. With continuous monitoring, threat modeling, and collaboration between security and development teams, organizations can stay ahead of emerging threats and ensure the long-term security of AI systems.

Wiz protects Microsoft Azure workloads:

- Get in-depth defense, including prevention, agentless visibility, and risk reduction
- Prioritize network and identity misconfigurations with a graph-based network and identity engine
- Get active detection and response with real-time monitoring
- Discover misconfigurations that compromise high-value assets

Wiz secures Microsoft Azure environments:

- Protect Microsoft Azure clouds in minutes with a 100% API-based solution
- Cover the entire cloud stack, including all VMs, containers, serverless, and PaaS
- Scan every Microsoft Azure layer without agents to detect vulnerabilities
- Model effective security posture by compiling all settings, compensating controls, and relationships





Wiz and Microsoft: Secure innovation with AI starts here

Wiz and Microsoft want to simplify cloud and AI security for customers. Through built-in security controls, data governance tools, and an emphasis on responsible AI principles, organizations can confidently use AI while minimizing current and emerging risks.

The future of AI does not need to be daunting. Wiz and Microsoft provide an unmatched combination of security and AI innovation — helping organizations see and secure everything they build in the cloud.

See how Wiz and Microsoft help organizations securely use and manage AI

