

2025 SURVEY

SANS 2025 AI Survey: Measuring AI's Impact on Security Three Years Later

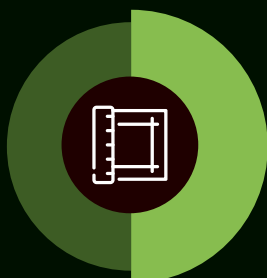
Written by **Ahmed Abugharbia**

September 2025

Key Findings

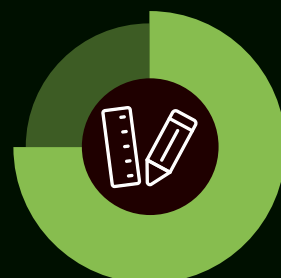
Security Is Lagging Behind Advancements in the AI Industry

Only half of the respondents who say their organization is currently using AI are using it for cybersecurity tasks.



Incident Response Teams Are Heavily Pursuing AI

75% believe AI will complement existing tools like SIEM, SOAR, and EDR over the next three years.



Robust Implementation Is Limited

Only **33%** use AI for investigating incidents.



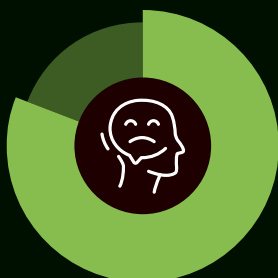
False Positives Are Overwhelming Analysts

The majority want to see fewer false positives in their reports and alerts. **66%** report that AI systems/agents generate many false positives, leading to alert fatigue.



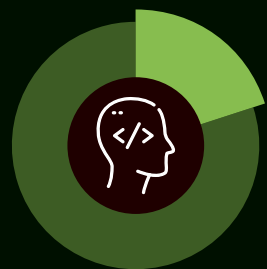
AI Attacks Are Feared More Than Defenses Are Used

Although AI is sparingly adopted by security teams, **81%** are concerned about emerging AI-powered threats.



Security Teams Are Not Involved Enough

Only 20% of respondents have limited involvement in governing generative AI (GenAI).

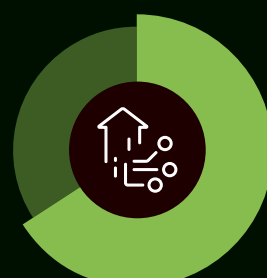


More Training Is Needed

51% say AI has affected security team training; **65%** emphasized the need for more specialized AI/cybersecurity courses; **64%** stressed the importance of continuous learning.

Respondents Are Optimistic AI Will Not Take Their Jobs

67% anticipate growing demand for professionals with AI and cybersecurity expertise in the next three years.



Survey Authors



Ahmed Abugharbia
SANS Certified Instructor

CURRENTLY TEACHING

SEC540: Cloud Native Security and DevSecOps Automation™

SEC545: GenAI and LLM Application Security™

SEC560: Enterprise Penetration Testing™

[VIEW PROFILE](#)

Ahmed Abugharbia is the co-founder of Cystack, a leading cybersecurity firm, and the founder of Cyberdojo, a company developing AI-based threat hunting tools and offering specialized cloud and AI security services. He has built advanced capabilities in cloud security and helped establish cloud security operations centers for multiple Fortune 100 companies, significantly strengthening their ability to detect and respond to evolving threats. With over 17 years of experience in the field, Ahmed has led large-scale initiatives across cloud security, network and application security, and incident response. A SANS Certified Instructor, he is the author of *SEC545: GenAI and LLM Application Security* and also teaches *SEC540: Cloud Native Security and DevSecOps Automation*.



Brandon Evans
SANS Senior Instructor

CURRENTLY TEACHING

SEC510: Cloud Security Engineering and Controls™

SEC540: Cloud Security and DevSecOps Automation™

[VIEW PROFILE](#)

Brandon is the owner and an InfoSec Consultant at On-Brand Technologies LLC, a consultancy helping organizations secure their applications and other workloads in multicloud environments, specializing in AWS, Azure, and Google Cloud. Prior to starting his consultancy, Brandon led the secure development training program at Zoom Video Communications. He began his career as a Software Engineer, where he worked on both the core product of a startup, later acquired by a Fortune 500 organization, and on various products spanning a multibillion dollar enterprise. Brandon is lead author for *SEC510: Cloud Security Engineering and Controls*, a contributor to *SEC540: Cloud Security and DevSecOps Automation*, host of *Cloud Ace podcast, Season 1*, an analyst for the SANS Multicloud Survey, a multi-year RSA Conference presenter, and recent Microsoft Defender Bug Bounty collector.

Introduction

AI is not on the horizon. It is here. It has been here. Generative AI (GenAI) and large language models (LLM) entered the cybersecurity zeitgeist nearly three years ago when ChatGPT became available to the public. Business leaders are clearly interested in using GenAI. They are scrambling to incorporate it in any way that makes sense—and many ways that might not make sense.

Security leaders are also interested in implementing AI. Similarly to the business as a whole, even if they do not have a fully fleshed out use case for GenAI, they do not want to fall behind their peers. Half of respondents stated that they are currently leveraging GenAI for security while 30% said they are planning to within the next 12 months.

How are early adopters applying GenAI to security? What security problems has GenAI made more manageable, if any? What problems and threats does GenAI pose to security teams? How should we expect GenAI-driven security to evolve over time? This report answers these questions using the survey responses from SANS's vast community of security experts.

GenAI, like anything else, is just a tool. Let us explore whether we should keep it in our utility belt at the ready or leave it in the toolshed for niche operations.

“ Expert Corner

I agree completely with the majority of respondents that security teams are lagging behind when it comes to AI adoption. In my opinion, this is because leadership teams are sure that they need AI but are not usually able to clearly articulate what that means. This, coupled with the dominant fascination with LLMs (which are amazing!) tends to impact the variety of AI/ML solutions in the cybersecurity space. Honestly, this is the reason SEC595 exists. I created it to teach SOC teams and threat hunters how to leverage AI and ML to create real-world solutions today that far exceed the commercial offerings available ... and without ever using an LLM!



[VIEW PROFILE](#)

David Hoelzer

SANS Faculty Fellow and
COO at Enclave Forensics, Inc.

COURSE AUTHOR

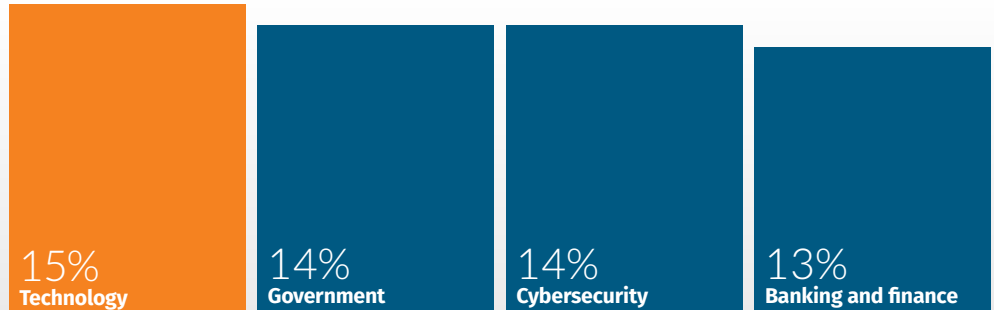
SEC495: Leveraging LLMs: Building & Securing RAG, Contextual RAG, and Agentic RAG™

SEC595: Applied Data Science and AI/Machine Learning for Cybersecurity Professionals™

Respondent Demographics

Most respondents were based in the United States (51%), with Europe second at 20%. The top industries represented were technology (15%), government (14%), and cybersecurity (14%) with the largest response from companies with fewer than 100 employees (18%). Figure 1 shows the survey demographics in detail.

Top 4 Industries Represented



Regions



Top 4 Roles Represented

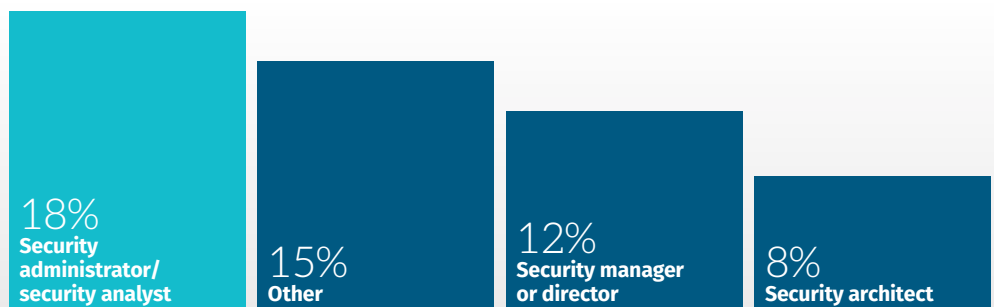


Figure 1. Demographics

Use Cases for AI in Security

Artificial intelligence continues to reshape cybersecurity. Its adoption varies across different cybersecurity disciplines. Although many security leaders see AI as a powerful tool they can use, the reality is more nuanced. Some domains, like application security, are already seeing meaningful integration and benefit. Others, such as incident response and red teaming, remain in earlier stages of adoption. In the following sections, we examine how organizations are navigating AI adoption in these three critical security areas.

AI in Incident Response

According to the respondents, incident response teams are currently adopting AI to a moderate degree. Only 26% of organizations use AI for responding to incidents and 33% for investigating incidents (see Figure 2). At the same time, 55% of organizations plan to incorporate AI into incident response for automated threat detection and analysis. Although this indicates a desire to grow in this area, it is unclear what the time frame for this growth would be. A possible reason for the low utilization of AI in incident response could be the maturity level of tools available. Outside of this survey, the authors have noticed a surge in the number of AI startups that focus on incident response over the past two years. This could mean that more advanced AI capabilities might be available soon.

Incident response teams are heavily pursuing AI.

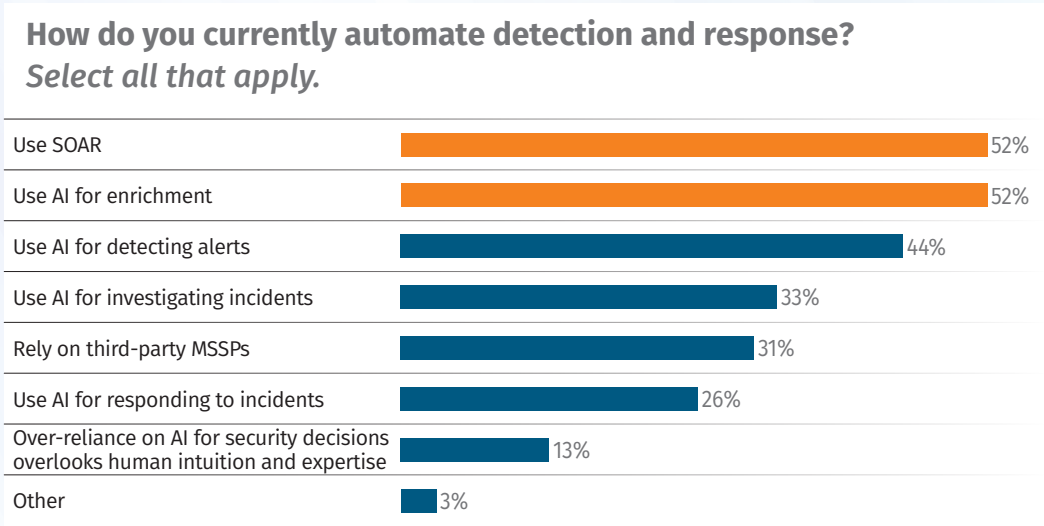


Figure 2. How Organizations Currently Automate Detection and Response

The primary applications currently in use are focused on supporting functions rather than autonomous action. Just over half (52%) use AI for alert enrichment. Although this can provide crucial context to security alerts, AI could potentially have a bigger impact in investigation and direct response. A potential reason for this lack of adoption could be the maturity of existing AI solutions.

Most of the survey respondents expect AI to be complementary rather than disruptive, with 75% seeing AI as a complement to existing tools like SIEM, SOAR, and EDR over the next three years and only 13% expecting complete replacement. This suggests the industry sees opportunities for AI to enhance operations rather than fundamentally shift in architecture. We will discuss our opinions about this assessment near the conclusion of this report.

Respondents also are concerned with significant AI challenges. 66% report AI systems generate many false positives, adding to the dreaded alert fatigue that incident responders already struggle with. Additionally, 58% cite heavy dependence on training data quality and 48% report AI struggles with context, leading to missed threats or incorrect prioritization.

The models we currently have are trained on diverse datasets that are not specific to cybersecurity. The authors believe that as vendors and startups build more AI tools for security, they also might start training or fine-tuning their own models for cybersecurity use cases. We expect that this also can address the context issue to some extent, and that these tools might be able to learn organization specifics such as employees and technology stacks. This also can contribute to understanding the context when investigating incidents.

“ Expert Corner

When two-thirds of teams report AI-driven noise, yet over half plan to expand automation, it's clear we're confusing urgency with readiness. False positives are not a glitch. They result from skipping the hard work of data curation, process integration, and precision tuning. To succeed, strategy must come before scale.



[VIEW PROFILE](#)

Seth Misenar

SANS Faculty Fellow

COURSE AUTHOR

LDR414: SANS Training Program for CISSP® Certification™

SEC511: Cybersecurity Engineering: Advanced Threat Detection and Monitoring™

These limitations highlight a critical tension: Incident response demands both speed and accuracy, but current AI systems often sacrifice one for the other. With 41% of respondents worrying about over-reliance on AI expertise, it suggests responders understand that incident response may require nuanced judgment that AI solutions have not mastered yet. For example, their responses could be enhanced with retrieval-augmented generation (RAG), especially when dealing with real-time data.

Despite these concerns, organizations are preparing for an AI-enhanced future by focusing on continuous learning (75% plan to promote AI education) and gradual integration.

AI in Application Security (AppSec) and Code Review

More than a third (37%) of organizations currently use AI in their AppSec activities, while 30% do not, and 32% are unsure. Just like incident response, the primary applications focus on security analysis enhanced with AI capabilities with the most augmented AppSec tool being static analysis security testing (SAST) at 65% (see Figure 3.)

Looking ahead, organizations are bullish on AI-driven code review capabilities. 71% anticipate AI agents leading or assisting code review will have the most significant impact within two years, while 57% expect that autonomous AI agents will review findings from other tools without human input. This suggests the industry sees code review as a prime candidate for AI automation.

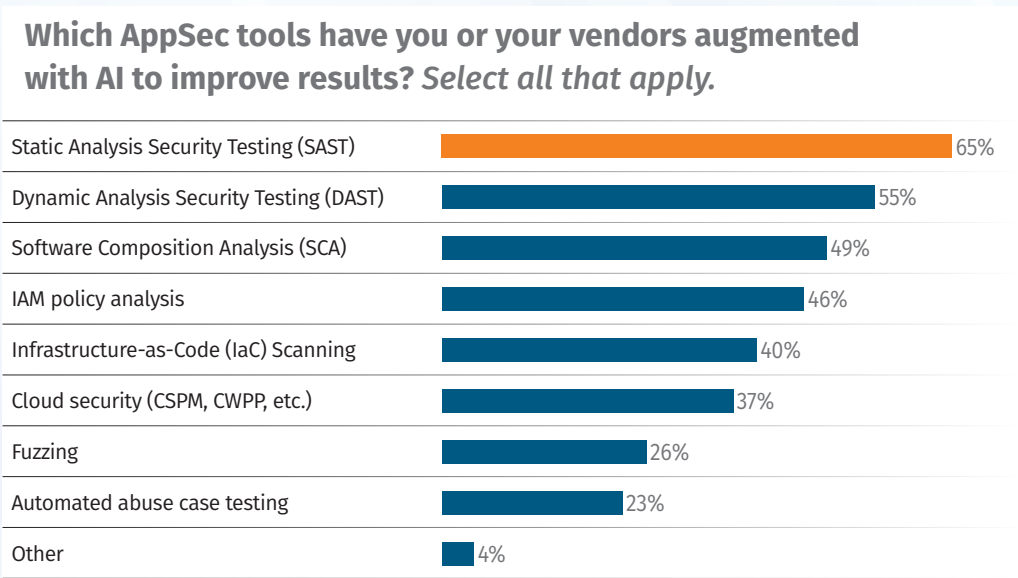


Figure 3. AI Usage for Security Tasks

However, challenges with AI and AppSec are similar to those seen in other security domains such as incident handling. Nearly 60% struggle with integrating AI tools into existing AppSec workflows and pipelines, while 51% cite high complexity and resource demands as a key challenge. The most concerning result is that 49% report issues with AI model reliability and potential biases impacting effectiveness.

AppSec false positives can delay shipping releases, while false negatives can ship vulnerabilities to production or delay detecting and fixing existing vulnerabilities. 44% of AppSec teams worry about keeping up to date with the rapidly evolving AI technologies. For these reasons, security organizations should be very intentional with the implementation of GenAI into their existing AppSec workflows.

AI in Red Teaming

Not many organizations are using GenAI for red teaming (15%). This aligns with interest respondents ranked “red team activities” as, by far, the least area they are planning to use AI technology in. This is somewhat surprising. GenAI can help create malware, phishing campaigns, and authorization bypass payloads that are tailored to the target organization. The respondents agreed with the tool’s potential: 53% of those using AI for red team activities think they can enhance collaboration and knowledge sharing between the red team and blue team by creating more realistic simulations.

Given how powerful these capabilities are, why might respondents be hesitant to adopt them? The main reason is ethics. In a normal red team campaign, security professionals should prove vulnerabilities exist by retrieving as little sensitive data as possible. Can we trust that GenAI-powered automation will apply the same level of care and consideration? The respondents do not believe so. The top ethical concern respondents have regarding GenAI and red teaming is whether it can adequately respect the privacy of real users (37%).

Despite GenAI’s potential to enhance red team operations with realistic simulations and tailored attack payloads, ethical concerns such as user privacy, bias, and unintended harm, are the top reasons security professionals are hesitant to fully adopt it for offensive security tasks.

There are several other ethical concerns including that GenAI's bias could skew results and hurt fairness (31%). Another quarter of respondents (24%) prioritized minimizing unintended harm to systems and data. This corresponds with the top challenge red teams have had incorporating GenAI so far: Avoiding automated attacks from causing real damage in a production environment (56%). Only 6% of respondents were most concerned about getting informed consent before launching AI exercises that impact others. This could be because, although red teams care about consent, real attackers do not.

At the same time, organizations are looking at how to use GenAI in the non-destructive portions of the red teaming process. More than half (59%) see value in using GenAI to summarize and format findings, while 55% would like to use GenAI to generate action items for these findings. Red teams should consider getting their feet wet with these activities before unleashing GenAI's full potential.

Using AI for Security Challenges and Concerns

AI implementation in cybersecurity reveals a significant gap between ambition and execution. While 50% of organizations currently use AI as part of their cybersecurity strategy and another 30% plan to start within 12 months, the depth of implementation remains shallow across most use cases.

The implementation pattern shows heavy concentration in a few areas with significant gaps elsewhere. For example, 53% focus on anomaly detection and 49% on alert enrichment, but adoption drops sharply for more complex applications such as using AI for code fix generation (18%) and red team activities (15%), even though these are areas where AI could provide substantial value.

When asked about the most significant challenges or limitations they face incorporating AI into AppSec efforts, 60% cited integrating AI tools with existing AppSec methods and 51% mentioned high complexity and resource demands. This suggests that more mature tools and more expertise in AI among cybersecurity professionals are needed.

Although AI adoption in cybersecurity is growing, most organizations remain stuck at the surface focusing on easier use cases, such as anomaly detection due to integration challenges and resource constraints that hinder deeper, high-impact applications.

AI Attacks Are Feared More Than Defenses Are Used

Security teams are reckoning with the fact that their adversaries also have access to GenAI platforms. While roughly half of the respondents' security teams are currently using GenAI, a whopping 81% are concerned with AI-powered threats. This implies that they believe attackers are more competent at GenAI than they are. It is interesting that defenders feel this way. Likely, attackers are also dealing with growing pains. Still, defenders should assume the worst to stay ahead of the most advanced threats.

Specifically, 83% are worried about highly personalized social engineering attacks. A related technique, deepfakes, greatly concern 73%. Additional fears include AI accelerating vulnerability discovery (67%) and helping to evade detection (59%).

Respondents are also highly concerned with attacks on the AI platforms themselves. 71% of respondents worry employees will pass sensitive data to GenAI platforms such as ChatGPT that could leak to other platform users, and 52% believe attackers can manipulate training data to provide their target with detrimental prompt responses. This emphasizes the need for security teams to provide standardized, enterprise AI tools for their employees to use instead of allowing them to pick their poison.

“ Expert Corner

The 2025 AI Survey has revealed a concerning disconnect that I've witnessed firsthand between red and blue teams. While 81% of security teams are concerned over growing AI-powered attacks, only 50% use these same tools for cybersecurity defense, and worse yet, just 33% leverage it for incident response. This gap between threat perception and defensive implementation suggests we're preparing for yesterday's war while tomorrow's adversaries are already weaponizing these capabilities. I believe the key to this disparity can be found in the fact that 66% of participants report that AI systems generate excessive false positives. Organizations are deploying AI without the necessary customization, context, and integration required for meaningful security outcomes. As security professionals, we don't deploy any other defensive solutions without first adapting them to our specific environments, so why is AI any different? Security leaders need to move beyond treating AI as a plug-and-play solution and instead invest in the foundational work of data quality, model tuning, and workflow integration. Otherwise, we risk creating security theater with expensive tools that increase analyst fatigue rather than enhancing our defensive posture.



Foster Nethercott
Certified Instructor Candidate

COURSE AUTHOR

SEC535: Offensive AI – Attack Tools and Techniques™

[VIEW PROFILE](#)

Governance and Ethical Oversight

AI governance in cybersecurity reveals a concerning gap between recognition and implementation:

- Most of the cybersecurity professionals surveyed (68%) believe they should have a role in governing AI use across their enterprises, though actual governance maturity lags significantly.
- Only 35% have a formal AI risk management and compliance program in place, while 42% are still in the early stages of developing policies.

This suggests many organizations understand governance is important but have not yet built the frameworks to manage it effectively.

Drivers for governance are split between external pressure and internal initiative, with 35% citing regulatory and legal requirements as their main driver and 39% pointing to internal risk management initiatives. Interestingly, only 16% are motivated by stakeholder concerns about AI ethics and bias and 10% do not consider AI auditability a priority. This pattern suggests compliance-driven rather than values-driven governance approaches may dominate.

It was also noticeable that third-party AI risk management shows more maturity than internal governance. Although 57% conduct risk assessments against AI-specific vendors, 41% do not have AI-specific controls. More concerning, 24% do not evaluate AI risks from third-party vendors at all, which is a significant blind spot considering how many organizations rely on AI-enhanced security tools from external providers.

The cybersecurity team's governance role centers on policy development rather than technical oversight. Nearly three-quarters (70%) participate in enterprise-wide AI governance policy development, while 40% establish incident response procedures for AI systems. However, only 23% conduct thorough testing for adversarial attacks, suggesting governance may focus more on process than actual security validation of AI systems.

In addition, lack of visibility into AI model use cases and risk exposure (56%) emerged as their biggest audit challenge, while 52% say they struggle due to a lack of established frameworks for AI risk assessments. When organizations cannot see what AI platform they are using or how to assess its risks, governance becomes very challenging.

Most cybersecurity teams see the need for AI governance, but poor visibility, limited risk controls, and a focus on policy over technical validation reveal a major gap between intent and implementation.

The data reveals organizations are caught between recognizing AI governance as critical and implementing it effectively. The emphasis on policy development over technical validation—combined with poor visibility into AI usage—suggests that the organizations are facing significant difficulties applying security controls on AI tools.

Bridging the Gap Between AI Governance and Security Practice: Mapping Survey Findings to SANS Critical AI Security Guidelines

The survey data shows five issues organizations face when implementing AI governance in cybersecurity. Although many organizations acknowledge that AI governance is critical, most seem to struggle to operationalize it effectively. The Critical AI Security Controls Guidelines is a framework that provides expert insights into helping organizations secure AI deployments, address evolving threats, and align security with scalability and governance needs.¹ See Table 1 for how these guidelines can assist with these issues.

Table 1. Recommendations from the Critical AI Security Controls Guidelines	
ISSUES	RECOMMENDATION
Organizations understand the importance of AI governance but lack operational structure. Most organizations believe cybersecurity should play a role in AI oversight, but few have implemented formal programs or governance to manage AI risk.	Governance, risk, and compliance (GRC) —Build an AI-specific governance foundation through mechanisms like an AI GRC board or extending current enterprise risk governance to include AI initiatives. Governance should not be passive. Organizations should define usage policies, track adherence to AI-related regulations, and build lines of accountability.
AI policies are created, but technical validation is missing. Although policy development is progressing, most organizations have not paired these policies with active technical testing or validation of AI system security and behavior.	GRC and monitoring —Governance must be reinforced with operational oversight and testing. Organizations should integrate adversarial testing and model red teaming into their development life cycle with regular reassessment of deployed models. Prompt and output logging, drift monitoring, and anomaly detection are all core recommendations. This helps ensure that models remain aligned with expectations and policy mandates over time, closing the gap between high-level governance and on-the-ground control.
External AI risk management outpaces internal efforts. Organizations often have stronger processes for evaluating external AI tools than for managing their own internal AI usage—and in some cases, they fail to assess vendor risks entirely.	GRC and deployment strategy —Emphasize treating vendor AI tools with the same scrutiny as internal systems. This includes validating the provenance of vendor-supplied models. By requesting AIBOMs from vendors, security teams gain visibility into the components and potential risks embedded in third-party models.
AI systems often lack fine-grained access control mechanisms. Access to inference endpoints, vector databases, and other AI related systems is often under secured, increasing the likelihood of tampering, leakage, or misuse.	Access control and data protection —Apply zero trust and least privilege principles across all AI components, including APIs, vector databases, and function-calling features. In RAG architectures, vector stores should be encrypted, access-controlled, and continuously monitored for integrity. Validate and sanitize all training and augmentation data before use. For agentic systems, explicitly scope and restrict function calls to prevent overreach. These controls extend proven access management practices to AI infrastructure, helping minimize the attack surface.
Organizations lack visibility into AI usage and risk exposure. Many organizations cannot see what AI models are being used within their infrastructure, where they’re deployed, or how to assess their risks—making governance efforts ineffective.	Monitoring —Visibility is foundational to the control set. Organizations must build continuous monitoring that tracks inference behavior, prompt content, usage volume, and other events. Internal and external AI systems should be treated like critical applications, with telemetry integrated into enterprise monitoring tools. Emphasis is also placed on protecting audit logs, as they may contain sensitive data. Without monitoring, other controls cannot function effectively—visibility is the prerequisite for actionable governance.

¹ “Critical AI Security Guidelines,” www.sans.org/mlp/critical-ai-security-guidelines

Workforce Impacts and Future Trends

The training transformation is comprehensive and demanding, with 65% reporting that it has required more AI-specialized training for cybersecurity and 64% emphasizing the need for continuous learning to keep up with the rapidly evolving AI technologies. The numbers reflect an understanding among the surveyed cybersecurity professionals that there is a gap in knowledge when it comes to AI technologies.

The impact of AI on the cybersecurity workforce is already substantial. 51% of organizations report AI has affected training requirements for their security teams, and 54% have observed job-related changes due to AI integration (see Figure 4).

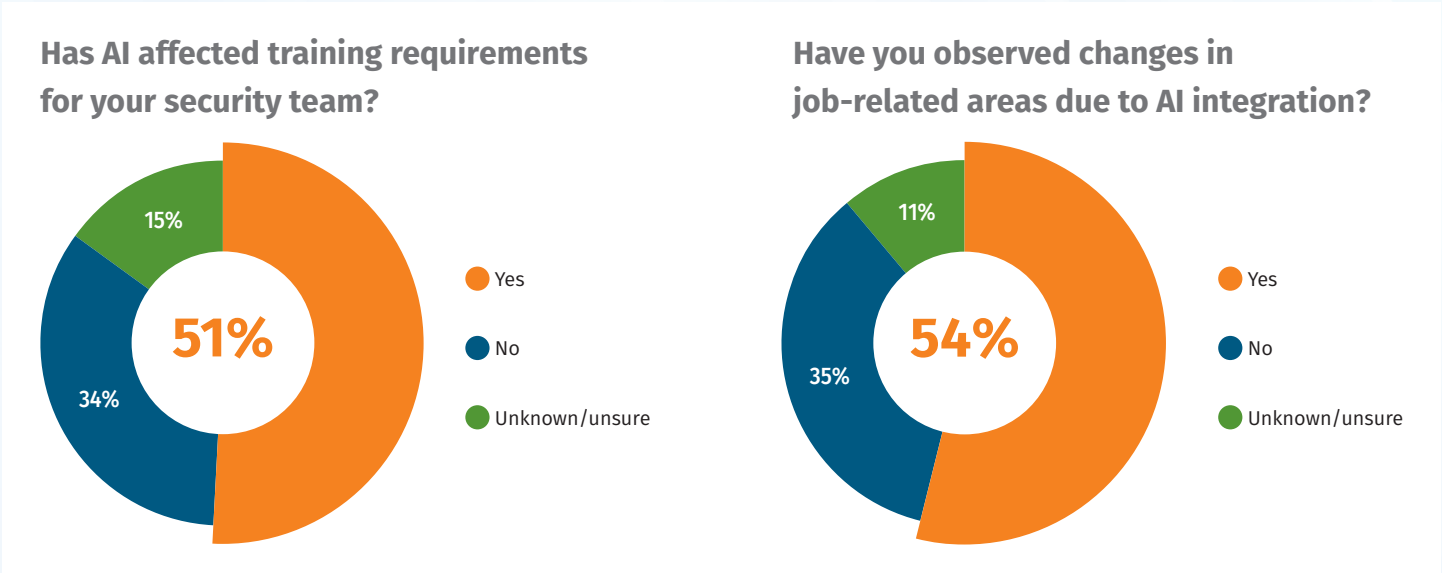


Figure 4. AI Effect on Training Requirements and Jobs

Looking ahead, organizations recognize this transformation requires investment. 58% have initiatives to prepare their workforce for the AI-driven landscape, mainly through providing ongoing AI fundamentals training (71%) and organizing workshops and hackathons (49%) (see Figure 5).

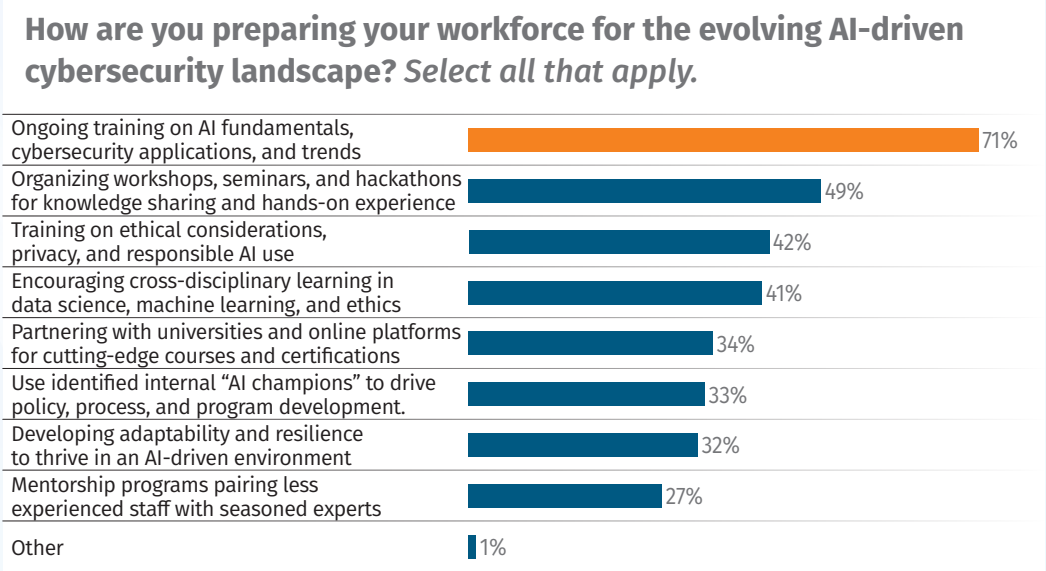


Figure 5. How Organizations Are Preparing Their Workforce for the Evolving AI-Driven Cybersecurity Landscape

AI Is Transforming Security Roles

We asked respondents to highlight the top misconception about AI in cybersecurity that they would like to dispel. The most popular answer (33%) was that “AI automates tasks but doesn’t replace human jobs; it shifts roles.” 11% also stated that, “AI processes data well but may lack nuanced context without human input.” As a result, 67% anticipate growing demand for professionals with AI and cybersecurity expertise in the next three years.

Frankly, we believe this is largely cognitive dissonance. A phrase that is frequently said by optimists is that “AI will not replace your job, but someone using it will.” Although this is technically true, it does not account for scale. As GenAI evolves, it is conceivable that five jobs will be replaced by a single professional using AI. As such, we unfortunately expect that the market for entry- and mid-level security professionals will contract in the coming years.

AI is already credited for eliminating security jobs. One example is the CrowdStrike layoff of 500 employees in May 2025.² SANS Institute Chief of Research Rob T. Lee argues that this is a huge mistake, comparing it to “cutting the fire department during wildfire season.”³ Still, business decision makers may feel the need to take that risk, especially in times of economic uncertainty.

However, this challenge poses security professionals with a fantastic opportunity. There are some security tasks where simple prompt engineering just will not cut it. For those cases, highly skilled and highly trained security professionals with GenAI expertise will be in high demand. Everyone is trying hard to keep up with this ever-changing field. The authors of this survey certainly are. We recommend that you take this opportunity to get ahead of the curve.

Respondents remain optimistic that AI won’t replace their jobs, but they acknowledge that it is fundamentally reshaping the nature of their work through automation of tedious tasks and advancement of skill development.

² “InfoSec Layoffs Aren’t the Bargain That Boards May Think,” May 2025, www.darkreading.com/cyber-risk/infosec-layoffs-arent-bargain-boards-may-think

³ “InfoSec Layoffs Aren’t the Bargain That Boards May Think,” May 2025

Conclusion

The survey data shows that the cybersecurity industry underestimates the transformational nature of AI. Organizations have made many first steps toward AI adoption. However, much more is needed.

This suggests the security community has not yet grasped the full scope of the AI transformation.

Security teams are implementing AI in basic use cases like alert enrichment and anomaly detection. They are not focusing nearly as much on transformational applications like autonomous code review, threat hunting, and incident investigation. Most respondents expect AI to merely “complement” existing tools rather than fundamentally reshaping security operations, which indicates that the industry is preparing for an incremental change while AI technologies are moving toward fundamentally changing computing.

This conservative mindset becomes even more concerning when viewed against the threat landscape. Although defenders are implementing AI in limited use cases, adversaries are not. They do not share the defenders’ ethical constraints, governance concerns, or technical limitations—they are already weaponizing AI capabilities. The AI knowledge gap is urgent, but the good news is that the industry recognizes its AI-knowledge deficit.

The authors’ assessment about the changes to cybersecurity jobs—that AI won’t replace jobs but will enable one professional to do the work of five—is not a pessimistic view. It is a likely potential outcome for any transformative technology. Security professionals who recognize this transformation early and adapt by developing AI expertise will thrive. Other types of cybersecurity roles also might emerge.

However, the window for getting ahead of this curve is narrowing. The organizations and professionals who treat AI as a fundamental shift will have a role in defining the changes in the cybersecurity field. Those who do not embrace it responsibly will risk irrelevance.

Many security teams still do not see AI as the transformative force that is it—leaving them unprepared for rapidly advancing threats and missing the opportunity to shape the future of cybersecurity.

Sponsor

SANS would like to thank this survey's sponsor:



About the SANS Research Program

The SANS Research Program is a key initiative by the SANS Institute and a premier global provider of cybersecurity research and information. SANS Research Program is designed to provide cybersecurity practitioners and leaders with data-driven insights, thought leadership, and solutions that help them better understand and respond to evolving security challenges. All content is authored by SANS instructor experts from around the world who apply their years of experience from hands-on practitioner work in the field, advisory roles, and the classroom to provide education, guidance, and actionable insights that help make the cyber world a safer place.

To learn about sponsorship opportunities for research, content, and in-person or virtual events, email us at **Sponsorships@sans.org** or go to **www.sans.org/sponsorship**.