

2026 SANS AI Survey Insights

Poisoned Wells and
Pure Springs: Drawing
Security and Compromise
from the Same AI Source

Written by
Matt Bromiley

July 2026



BACKSLASH



FORTINET



infoblox



XBOW



About Our Authors

AI Survey Insights Author

Matt Bromiley, Certified Instructor

Matt Bromiley is currently with Prophet Security and has 12-plus years of experience in the cybersecurity industry making life difficult for cyber threat actors. He also serves as a GIAC Advisory Board member, a subject-matter expert for the SANS Security Awareness group, and a technical writer for the SANS Research Program, including lead author on the **Critical AI Security Guidelines**. Matt brings his passion for incident response and leadership to the classroom as a SANS Instructor for **LDR553: Cyber Incident Management**.



[▶ READ MORE ABOUT MATT](#)

AI Survey Designer

Foster Nethercott, Certified Instructor Candidate

Foster Nethercott is a SANS course author for **SEC535: Offensive AI – Attack Tools and Techniques** and a cybersecurity professional specializing in offensive operations, as well as AI-driven threat development, tools, and techniques. A US Marine Corps veteran with nearly a decade of experience in cybersecurity, Foster brings a mission-focused mindset to his work—translating real-world offensive tactics into practical skills that help students understand how modern attackers think, adapt, and evade detection.



[▶ READ MORE ABOUT FOSTER](#)

Cyber Leaders

Alongside the main survey, a separate group of senior security leaders (such as CISOs, CSOs, and security VPs) completed a dedicated survey module. Their answers show where leadership priorities line up (or don't!) with the practitioners' day to day.

Key Highlights

- 1** Leaders and practitioners are not describing the same governance program.

50%

of leaders report a formal AI risk management program; only **36%** of practitioners say the same. The questions use identical wording, so the 14-point gap is a real perception difference, not a wording artifact. A program that the people doing the work cannot see is not governing much in practice.

- 2** Leaders are leaning into AI for defense faster than into defending against it.

63%

of leaders have shifted priority toward using AI for defense over the past year; **16%** have shifted toward protecting against AI-enabled threats. Meanwhile **78%** of organizations in the practitioner sample already saw AI-enabled attacks. The defensive investment is sound. The exposure on the offensive side is arriving faster than the priority shift reflects.

- 3** Vendor accountability sits at the top of the leadership worry list.

42%

of leaders rank vendor efficacy and data responsibility as their joint top concerns, each cited by **42%** of respondents. The skills gap and the transparency problem tie for third at **33%** each. The leadership view is pointed outward, at the AI supply chain, more than at the tools' day-to-day reliability.

- 4** Stated governance maturity runs ahead of what exists on the ground.

50%

of leaders report a formal program, yet **44%** also describe their organization as still in the early stages of writing AI governance policy. Some are claiming both at once, which suggests "formal program" is carrying a lot of weight. The commitment is genuine, and most leaders are still working out how to turn it into something with real operational reach.

Where Leaders Are

Seventy-six percent of leaders say their organization is actively using AI for cybersecurity, almost identical to the 78% in the broader practitioner population. Deployment depth is where the two groups differ. Among leaders using AI, 44% call their deployment early production and 25% are still piloting. Twenty-eight percent have reached mature production, and a single respondent described AI as mission critical. That last figure sits well below the 11% of practitioners who say the same, which points to teams operating AI deeper in production than the leadership view registers.

The Priority Shift

Leaders answered one question their teams did not: Over the past year, has your priority moved toward using AI for defense or toward protecting against AI-enabled threats? Sixty-three percent moved toward defense (22% significantly, 41% somewhat) and 16% moved toward threat protection. The rest reported no change or said they do not separate the two.

95% of leaders believe attackers are already using AI. Only 16% have shifted priority to defend against it.

Favoring defense is reasonable. AI gives real leverage in detection, investigation, and response, and building that capability is a defensible call. The complication is in the same survey, 78% of organizations already saw confirmed or suspected AI-enabled attacks, and 95% of all respondents believe threat actors are using AI, 58% of them significantly. Prioritizing defense is fine. However, organizations should not be treating adversarial AI as a horizon item.

Where Leaders Diverge from Practitioners

The governance perception gap is the sharpest divergence in the data. Fifty percent of leaders report a formal AI risk program; 36% of practitioners do. Either the programs exist mainly on paper, or the two groups mean different things by “formal.” Both readings point to the same conclusion: Governance must show up in the workflows, audit steps, and accountability that practitioners can name.

The two groups are also solving different problems. Leaders worry about the vendor relationship, namely efficacy and data handling, both at 42%. Practitioners worry about the tools in their hands, naming transparency (40%) and efficacy of AI provided by commercial vendors (38%) first. Both are facets of the same trust problem.

What Cyber Leaders Need to Act On

Three priorities come straight out of the leader data, each sitting on a gap between what leaders report and what practitioners live with.

1: Operationalize governance, don't just declare it.

Half of leaders believe a formal program exists, while practitioners disagree. Closing that gap means putting governance into review steps, approvals, audit trails, and model-use visibility that someone on the team can point to and describe, not just policy language in a form.

2: Bring adversarial AI into the threat model now.

The 78% already facing AI-enabled attacks are the reason the defense-first framing needs a counterweight. Board conversations should treat adversarial AI as a present operating risk. As attackers use AI to speed up reconnaissance, impersonation, and data discovery, the damage a single compromised identity can do grows, which makes limiting access a necessary security control.

3: Treat the skills gap as an operational problem, not a hiring line.

A third of leaders name the skills gap a top challenge. Solving it mainly through future hiring leaves the people already in the seats without the AI literacy to check outputs, override bad guidance, and build governance that functions.

**Adversarial AI isn't on the horizon.
It's already inside 78% of organizations.**

2026 AI Survey Key Findings

The 2026 survey describes a field that made a sound bet on AI and is now finding out what it costs to keep, as we saw adoption rise *sharply* in the past year. The governance, validation, and workforce depth needed to support that scale did not keep up. Adversaries adopted AI as fast as defenders, sometimes faster.

78%

Adoption rose sharply, but most deployments are shallow.

Active AI use in cybersecurity strategy went from 50% in 2025 to 78% in 2026. This is a 28-point jump, the largest this survey has recorded. However, only 27% of practitioners call their deployment mature production. Most are still piloting or running AI in a supporting role.

63%

More deployment has produced more reported failures.

Sixty-three percent of practitioners report significant AI shortcomings in threat detection and response, up from 45% in 2025. Two-thirds say AI guidance has steered them wrong at least once in the past year.

40%

Trust and transparency have replaced integration as the top barrier.

In 2025, the hardest part was wiring AI into existing systems. In 2026, the leading concerns are transparency in AI decisions (40%) and efficacy of AI provided by commercial vendors (38%). Neither is fixable with a configuration change.

61%

AI in red teaming went from minority to majority practice in a year.

Sixty-one percent of practitioners now use AI in red team work, up from 33% in 2025. Yet again, the steepest single-year jump in the survey. The top operational worry is keeping automated attacks from doing real damage in production (52%).

78%

AI-enabled attacks are being observed, not just anticipated.

Seventy-eight percent of organizations report confirmed or suspected AI-enabled attacks in the past year, and 95% believe threat actors are using AI. The techniques are spread broadly across phishing, exploitation, deepfakes, and reconnaissance.

76%

Governance responsibility grew. Governance maturity did not.

Seventy-six percent of security teams now hold a governance role for enterprise AI, up from 68% in 2025, but the share with a formal AI risk program barely moved year over year. The mandate expanded while the infrastructure stalled.

73%

Training requirements jumped as AI spread across every domain.

Seventy-three percent of practitioners say AI changed their team's training requirements in 2026, up from 51% the year before.

The above Key Findings and the sections that follow cover the full practitioner survey, including results, year-over-year trends, and the analysis and recommendations that come out of them.

Introduction

The adoption of AI is no longer a question, a “maybe,” or a “we should.” It’s here. With 78% of practitioners now using AI as part of their security strategy, the adoption question has largely been settled. What remains open is harder: whether organizations built the infrastructure to use AI responsibly, whether they can govern it consistently, and whether they are ready for the version already enabling adversaries.

Three threads run through this year’s data:

1. The first is a widening distance between deployment and confidence. We see that practitioners are using AI in more places *and* reporting more failures at the same time.
2. The second is governance responsibility outpacing governance capability, as security teams inherit accountability for enterprise AI without the frameworks or staff to act on it.
3. The third is offensive acceleration, with AI now a primary capability in both red team operations and adversary playbooks.

These threads feed each other and will be the basis for our survey insights.¹

An organization that cannot reliably judge whether its AI is right cannot govern it well, ungoverned AI ends up deployed without guardrails, and the offensive pace on both sides raises the cost of every gap.

¹Survey demographics are located at the end of this report.

The State of AI Adoption

Active AI use in cybersecurity strategy rose from 50% in 2025 to 78% in 2026, the largest year-over-year move the survey has seen. The share of organizations with no plans to adopt fell from 9% to 4%. For most of the industry, the decision to use AI in security has been made; the remaining 18% are mostly planning to start rather than holding out.

Deployment depth is the more revealing number. Among the 78% using AI, a third (33%) call it early production, meaning deployed but not essential and another 21% are still experimenting (see Figure 1). Most practitioners sit between testing and light operational use, which means shallow deployment is producing a partial picture. Organizations see enough value to keep investing but not enough depth to surface the failure modes that only appear at scale, which can leave confidence sitting on a thin floor.

In 2025, a real share of the field was still arguing about whether AI belonged in security workflows. By 2026, that argument is settled. The open question is whether the operational scaffolding was built to match the deployment pace.

Q: How would you characterize your AI deployment maturity?

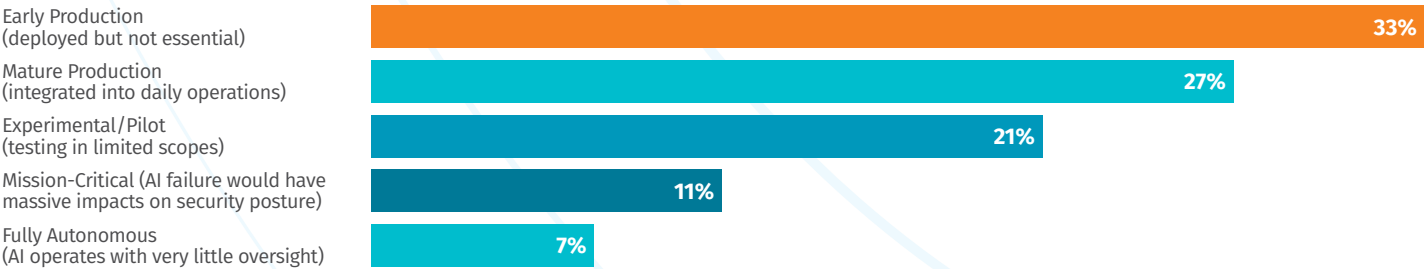


Figure 1. AI Deployment Maturity

AI is arriving through every channel at once with Microsoft Copilot in the lead (see Figure 2). Enterprise licenses, personal tools, and open source all show up together, which is itself a governance fact.

Whether that use is governed is a separate matter: Only 41% of organizations report using generative AI for security-related tasks under strict policy, while 39% say usage is informal with no policy at all. Another 13% want to use it but do not yet. The informal 39% represents a lot of AI activity running outside any organizational visibility, which connects directly to the governance discussion later in this report.

Q: What AI tools or platforms is your organization using for cybersecurity?

Select all that apply.

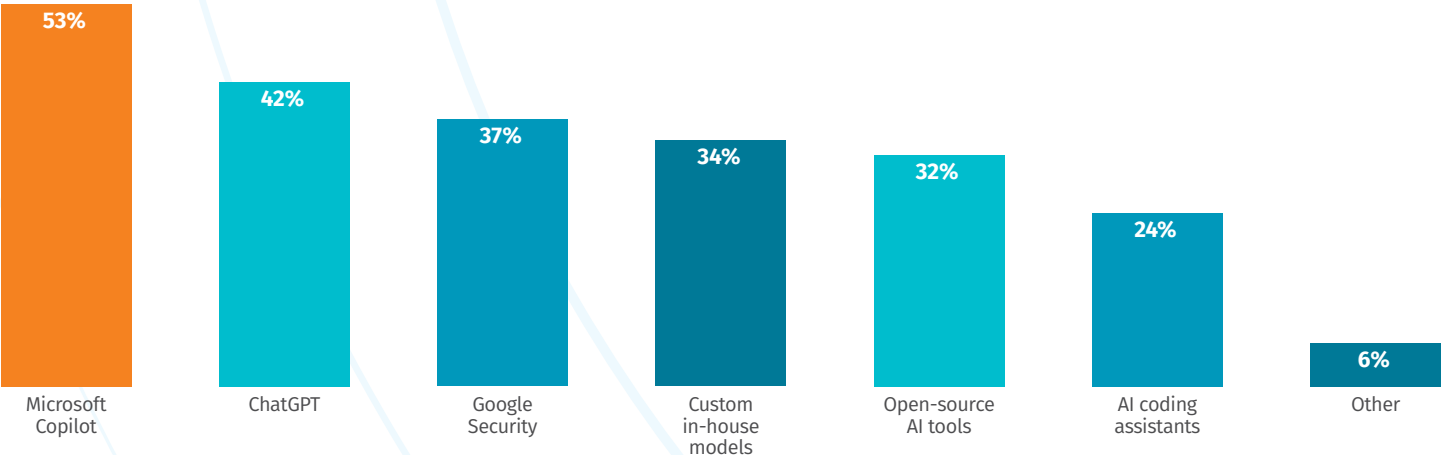


Figure 2. AI Tools and Platforms in Use

For specific security tasks, generative AI is already routine. Log analysis leads, followed by explaining threats, and writing code (see Figure 3), and the practitioners doing these tasks are not running pilots. The analysts using AI to triage logs are doing their existing jobs with a different tool set than they had a year ago. Whether they are doing it under policy is, again, the governance question.

Q: What security tasks do your teams use generative AI for?

Select all that apply.

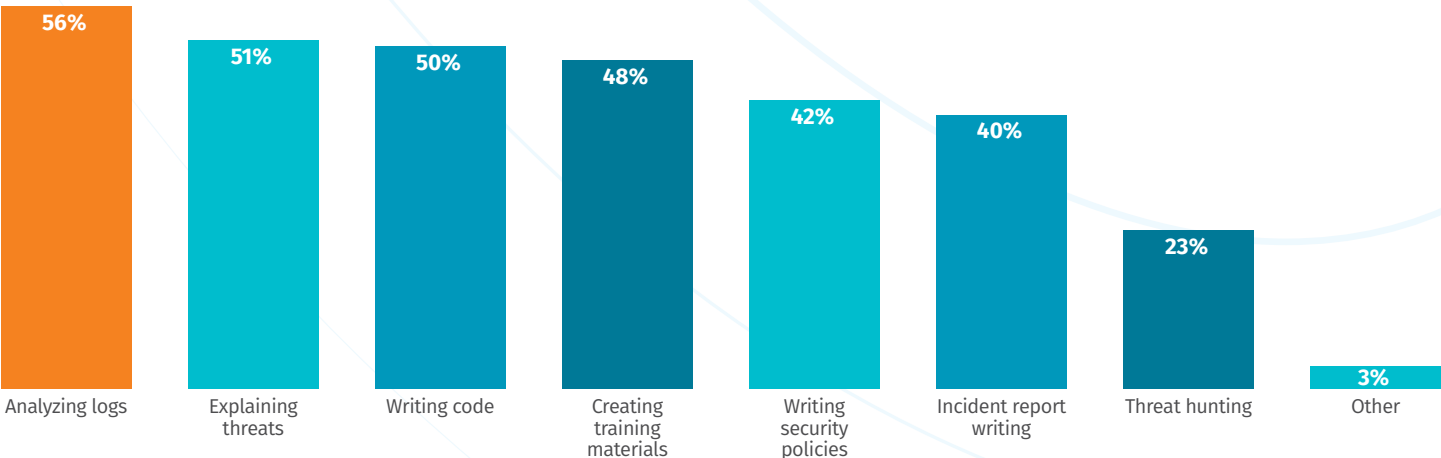


Figure 3. Top Generative AI Security Tasks

Where AI Is Working, and Where It Falls Short

The most common AI use cases in 2026 show where practitioners have decided to trust the technology. Incident investigation leads at 47%, then anomaly detection (39%), automated incident response (39%), forensic investigation (35%), and summarization of security issues (35%). A year earlier, anomaly detection led at 53% with incident investigation third. The shift is meaningful because it indicates that practitioners have moved AI further along the decision chain, out of pattern spotting, and into analysis and investigation.

Sixty-three percent of practitioners report significant shortcomings when AI detects or responds to threats, up from 45% in 2025. The failures cluster around false positives and alert fatigue, trouble recognizing genuinely new threats, and hallucination, representing confident, wrong output. For most teams running AI in production, these are the routine experience.

The misdirection data puts a number on it (see Figure 4). Asked how often AI guidance had pointed them the wrong way in the past year, 32% of practitioners said never. The other two-thirds had been misled at least once: 23% one to five times, 20% six to 10 times, 13% 11 to 20 times, and 9% more than 20. Teams that adopted AI to cut their workload have, in many cases, picked up a new error-checking workload alongside the old one.

“We tend to treat our AI as a digital intern and check its work. If we don’t exercise the capabilities ourselves, when the AI fails, we’ll lack the skills to address the security issues.”

—2026 SANS AI Survey respondent

Q: In the last 12 months, how many times has an AI security tool provided guidance that led your team in the wrong direction?

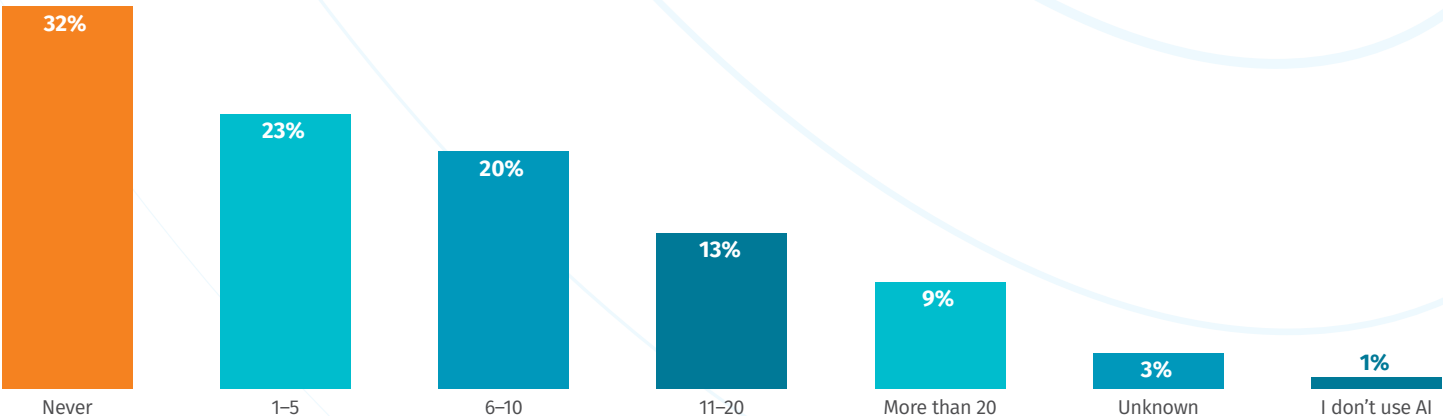


Figure 4. Times AI Guidance Led Teams in the Wrong Direction, Past 12 Months

Where practitioners want improvement is telling. The top requests are AI agents that adapt to threats in real time (51%), better integration of AI into security platforms (43%), and AI algorithms that reduce false positives (43%). The leading ask is for the capabilities *already deployed to work more reliably*. **Practitioners are not asking for more AI so much as for the AI to do its current job better.**

Does AI Actually Work? How Practitioners Measure It

Given how much AI is now deployed, the simplest question may be the most important: Is it working? The data says a qualified yes. Practitioners keep moving AI into more critical parts of the workflow, which only makes sense if it delivers enough value to justify the expansion despite the reported failures. The qualifier is that much of the confidence is hard to verify, because of how organizations measure success.

The metrics organizations track reveal where their attention sits. Time and cost savings from reduced manual work lead at 45%, followed by reduction in false-positive alerts (41%) and the quality of AI-generated threat intelligence (39%). These are efficiency measures, and they reflect what sold most teams on AI in the first place. Further down, 38% track detection-to-response time and 37% track precision.

What sits low on the list is what's concerning. Only 25% track recall rate, the share of real threats AI caught, and just 17% track false-positive rate as a formal metric. Efficiency is the easiest benefit to see and the easiest to overweight. An AI system can cut analyst workload sharply while still missing a meaningful fraction of genuine threats, and a team watching only efficiency may not notice until an incident forces the issue. The practitioners tracking precision and recall next to efficiency are the ones building an honest picture of performance.

Validation methods follow the same pattern. Fifty-two percent run periodic reviews, 49% do incident post-mortems, 41% run AI analyses in parallel with traditional systems, and 29% compare against industry benchmarks. All are legitimate, and all are labor-intensive. Periodic reviews carry a built-in limit: They check the system at a moment in time. AI behavior drifts as the threat landscape shifts, and a point-in-time review can miss that drift between checks. The 41% running parallel analysis are doing the most rigorous real-time comparison, but 41% is not a majority. For now, the field is spending human effort to cover AI's reliability gaps, which holds up in the short term but does not scale.

Expert Corner

“The pattern we see in the field, and in this data, is that organizations are adopting AI far faster than they're able to govern it. Closing that gap is exactly one of the intentions behind the AI Security Maturity Model. The most common mistake is treating a blanket block policy as risk management. In practice, telling people 'don't use AI' just drives usage into the shadows where management is even tougher. The goal is to pull it into visibility where it can be governed. Name an owner, inventory your AI tools, write a one-page policy, brief your team, and reassess in 90 days. That single sequence moves most organizations from Stage 1 to Stage 2 governance in a quarter, and a small company at a real Stage 2 is in a stronger position than an enterprise claiming a Stage 3 it can't prove.”



Chris Cochran

Field CISO and Vice President of AI Security
at the SANS Institute

Author of the SANS AI Security
Maturity Model eBook

[VIEW PROFILE](#)

The Trust Problem

The top integration challenge changed shape between 2025 and 2026. A year ago, it was wiring AI into existing systems, largely seen as an engineering problem. Now the challenges are transparency and trust in AI decisions (40%), judging vendor efficacy (38%), and hallucination (34%). The technical lift has gotten easier, but trusting the output did not.

The transparency problem has no quick fix and that kind of gap does not close with a patch. If a practitioner cannot tell why an AI system reached a conclusion, the only choices are to accept it on faith or to build a manual review process that eats away at the efficiency the deployment was supposed to deliver.

Confidence varies a lot by task. Practitioners are most willing to let AI act without review on threat classification and vulnerability prioritization, while they are least willing on behavioral anomaly detection and true-positive identification (see Figure 5). The split tracks where AI tends to do well, in structured classification, against where it struggles, in novel and context-dependent calls.

The low confidence in true positive identification is the one to watch, because confirming a real threat is the foundation of everything downstream in response.

In 2025, the hardest part of AI in security was getting it integrated. In 2026, it is knowing whether the system is right. Plugging AI in turned out to be tractable. Trusting its output has not.

Q: How confident are you in AI-generated security decisions without human interaction in the following domains?
Select all that apply.



Figure 5. Confidence in AI Decisions without Human Review, by Task

Expectations about AI replacing existing tooling have moved too. In 2025, 71% expected AI to complement SIEM, SOAR, and EDR, and only 13% expected replacement. By 2026, the complement view fell to 59% and the replacement expectation grew to 22%. Whether replacement arrives or not, a 9-point swing toward displacement in a single year is a useful signal of where practitioner confidence is heading.

AI as an Offensive Capability

Two numbers frame the offensive picture. Seventy-eight percent of organizations observed confirmed or suspected AI-enabled attacks in the past year (45% confirmed, 33% suspected). Ninety-five percent believe threat actors are using AI, 58% of them significantly. The adversary side is not a forecast anymore.

The practitioner side moved just as fast. Sixty-one percent now use AI in red team work, nearly double the 33% of 2025 and the steepest single-year jump in the survey. The benefits practitioners cite most are generating remediation steps and action items from findings (61%) and more realistic threat simulations (52%). AI is increasing both the volume and the quality of offensive testing.

Observed attack techniques are spread remarkably evenly (see Figure 6). Deepfake content edges ahead, with AI-assisted exploitation and AI-generated phishing close behind, and adversarial attacks on models, automated reconnaissance, and AI-powered brute forcing all clustered just under. The near-even distribution across techniques is itself the finding. Adversaries are integrating AI across the whole attack life cycle rather than concentrating on a single method. This means there is no single countermeasure that addresses the full range of what organizations are already facing. AI-enabled attacks also extend beyond faster phishing or automated exploitation. A compromised identity or a manipulated AI workflow can reach whatever sensitive data that identity or workflow is permitted to touch, so what an attacker can extract depends heavily on how tightly access is scoped.

Q: What AI-enabled attack techniques have you observed or experienced? Select all that apply.

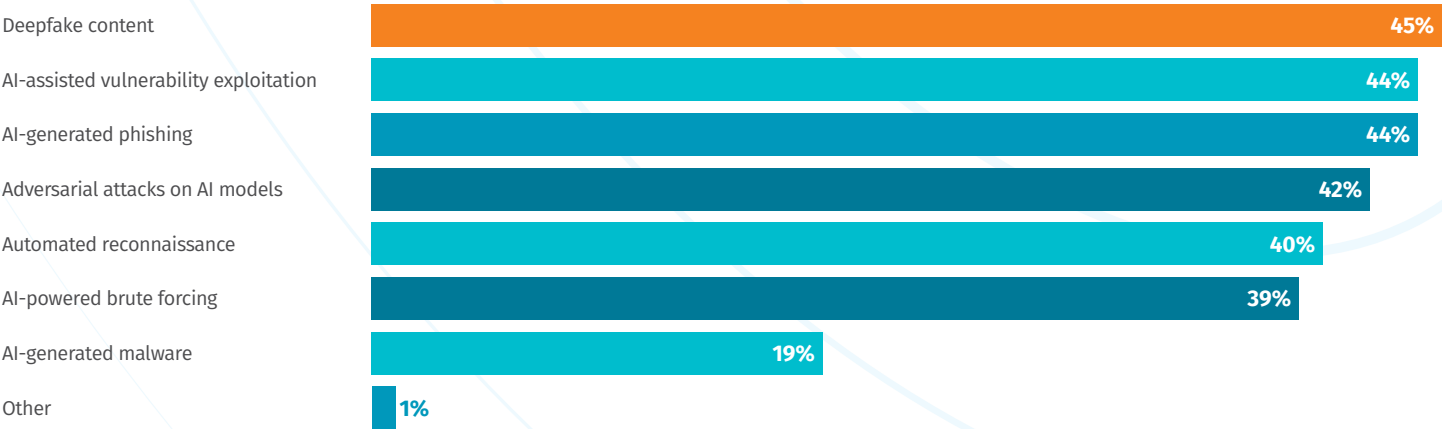


Figure 6. Observed AI-Enabled Attack Techniques

One thing worth noting is that stated concern about AI-powered phishing dropped from 83% in 2025 to 66% in 2026, and deepfake concern from 73% to 61%. The likely explanation is not that practitioners worry less, but that these threats moved from anticipated to operational. A team already countering AI phishing every week answers a “what worries you?” question differently than one bracing for it.

For red teams, the top challenge is not budget or tooling. It is keeping automated attacks from causing real damage in production, cited by 52%. A human tester hesitates before a destructive action. An automated process must be told to hesitate, in advance and explicitly. The 52% says this is a live operational problem, not a thought experiment.

Application security tracks the same curve. Sixty-seven percent of practitioners now use AI in AppSec, up from 37% in 2025. The tools most often augmented are DAST (51%), SCA (50%), SAST (43%), and IaC scanning (39%). The challenges echo red teaming almost exactly: complexity and resource demands (44%), ethical and legal concerns (42%), and model reliability (42%), with in-house expertise close behind at 41%. The same constraints showing up in both domains is the point. The skills gap and reliability concerns are not local problems with local fixes; they are structural, and they will cap AI’s effectiveness everywhere until they are addressed at the program level.

Practitioners expect the offensive impact to keep climbing. Fifty-two percent expect autonomous AI to find and exploit vulnerabilities faster than humans, and 45% expect AI attack agents operating without human input. If both land at scale, the gap between a vulnerability existing and being exploited compresses hard.

Expert Corner

“What stands out to me is the jump from 45% to 63% of practitioners reporting real shortcomings in AI threat detection and response, and I don’t read that as AI getting worse so much as teams finally running it at a scale where the cracks show. I’ve watched this pattern before with every major shift in security tooling. The technology outpaces the discipline needed to run it well, and for a while the gap just sits there quietly until enough people are relying on it to notice. AI can genuinely speed up detection and response, but only when the analyst behind it still has enough judgment to push back on what it hands them. The skill worth building isn’t learning to prompt the tool well. It’s knowing when to stop trusting what it gives you.”



Dave Shackelford

SANS Senior Instructor

COURSES TAUGHT

SEC501:

Advanced Security Essentials - Enterprise Defender (course co-author)

SEC504:

Hacker Tools, Techniques, and Incident Handling

[VIEW PROFILE](#)

What Is Working Against AI-Enabled Attacks

The more useful question than what attacks look like is what stops them. Asked which defenses have proven most effective against AI-driven threats, practitioners pointed mostly at behavioral and human-centered controls rather than AI-specific ones (see Figure 7).

Q: What defensive strategies have proven to be the most effective against AI threats?
Select your top 3.

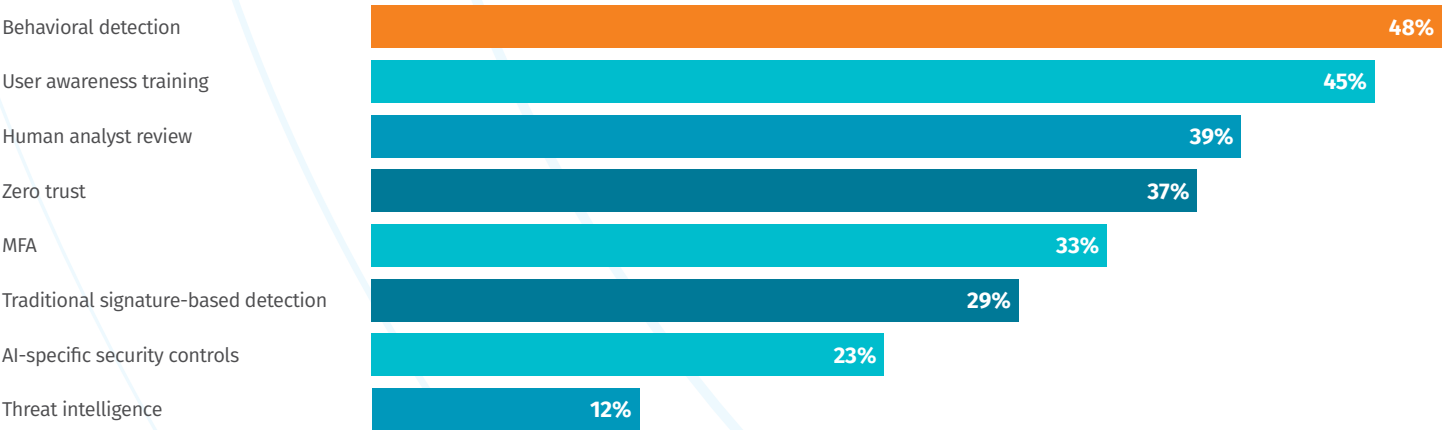


Figure 7. Most Effective Defenses Against AI-Enabled Threats

The pattern says defenders are not matching a threat category to a specialized countermeasure. Behavioral detection works against AI-enabled attacks for the same reason it works against any capable adversary: It watches what attackers do, not which tools they used. User awareness training landing second is worth sitting with because AI phishing and deepfakes both depend on fooling a person, which makes a harder-to-fool person a direct countermeasure. Human analyst review at 39% reinforces the same idea: A trained, skeptical human is the most durable defense against threats that are getting better at producing convincing output.

Asked for recommended practices, practitioners named strikingly foundational ones. Aligning AI initiatives with overall security strategy and business goals leads at 59%, then adopting transparent and explainable AI (52%) and rigorously testing models across scenarios before full deployment (52%). The top two map directly onto the largest failure modes in this survey. Organizations that deployed fast without working through alignment or explainability first are, by their own practitioners' account, paying for that sequencing now. The 52% calling for rigorous pre-deployment testing matters because most deployments are still early-stage, which means that testing window is, for many teams, still open.

Governing AI in the Enterprise

Security teams have taken on real governance responsibility for enterprise AI, with 76% now holding that role, up from 68% in 2025. In 2025, 42% described their AI risk management as early-stage and 35% had a formal program. In 2026, those figures are 43% and 36% (see Figure 8). The distribution barely moved in a year. The responsibility grew, but the foundation under it did not.

Q: How would you describe your organization’s current approach to AI risk management?

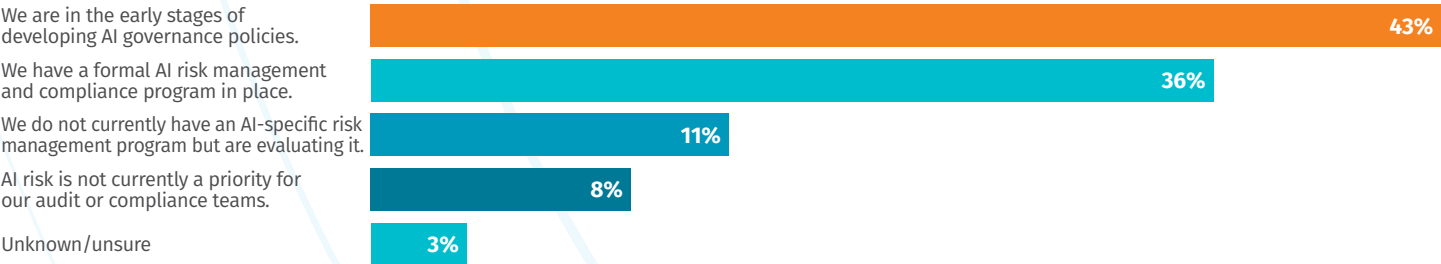


Figure 8. AI Risk Management Maturity (Practitioners)

The governance work practitioners do skews toward point-in-time checks. The most-reported activity is AI security assessments and penetration testing (30%), then collaborating with IT and data teams on secure deployment (26%), developing risk frameworks and policies (19%), and monitoring for breaches (16%). An assessment is a snapshot. It does not watch model behavior between evaluations, enforce policy in daily workflows, or catch output drift over time. The field is comfortable running an evaluation and thinner on the continuous oversight that governance requires.

Third-party AI risk shows the gap between intent and practice. Sixty-nine percent of organizations now run AI-specific third-party assessments, up from 57% in 2025. But 47% of that same pool also trusts vendors to manage their own AI risk without auditing them, up from 27% a year ago. Those two positions coexist inside the same organizations, which is only possible if the assessment is a form-filling exercise rather than a substantive probe. The question is not whether organizations assess third-party AI risk. It is whether the assessments test anything that matters.

Practitioners' biggest governance challenge backs this up. Sixty-three percent cite a lack of visibility into where AI models are used and what they expose, up from 56% in 2025. Fifty-four percent say there are no established frameworks for AI audits, essentially unchanged from 2025 (52%). These are infrastructure gaps, and another assessment form does not close them. Without visibility into what AI is doing and without shared standards to judge it, governance stays descriptive instead of prescriptive.

That visibility problem now extends to the data side of AI use. Thirty-six percent of practitioners already worry about sensitive data or company IP leaking through employee use of AI tools, and increasing visibility into where AI is being used is the single most common adaptation organizations report planning (52%). As copilots and AI agents inherit the permissions of the users and service accounts behind them, the data those systems can reach becomes a governance question in its own right: what regulated or proprietary information AI can retrieve, whose data it represents, and whether that access should continue. Scoping access tightly for both human and non-human identities limits what a compromised account, a manipulated workflow, or an over-permissioned agent can assemble and expose, which makes least privilege a practical AI governance control rather than a separate IAM concern.

Finally, the primary driver for AI governance is internal risk management (43%), ahead of board-level concern (24%) and regulatory requirements (24%). Building governance from within is a professional choice with real value. It also makes those programs easier to deprioritize when other demands crowd in, since internal initiatives lack the persistence that regulatory mandates carry.

Not every respondent felt they were able to effectively govern AI use in their organization, saying their role is to “Draft governance and provide advice for leadership to ignore.”

The Workforce Impact

The shortage of AI-literate security professionals is not a separate workforce topic sitting beside the operational findings. It runs underneath all of them. It is part of why misdirection goes uncaught, why trust in AI output lags deployment, why red teams struggle to contain automated testing, and why governance stalls at the policy-writing stage. Every operational gap in this survey has a workforce dimension, and none of them ease without people who can work with AI confidently enough to catch what it gets wrong.

The training impact has accelerated sharply. In 2025, 51% of practitioners said AI had changed their team's training requirements. In 2026, it is 73%, a 22-point jump in a year. AI has embedded itself across detection, response, red teaming, AppSec, and governance deeply enough that nearly three in four teams have rethought what their people need to know.

Sixty-eight percent of practitioners report AI-driven changes to their jobs, up from 54% in 2025. The most common shifts: traditional roles now carry AI oversight and integration duties (50%), continuous education is required to keep up (46%), and strategic work replacing routine tasks as AI absorbs the day-to-day (44%). These describe a reorientation of what a security practitioner is expected to do and know, not a minor adjustment.

Job satisfaction adds nuance. Sixty-six percent say AI has affected satisfaction, up from 52% in 2025, mostly for the better. The top response is a greater sense of accomplishment from making AI integration work, cited by 70% of those affected, up sharply from 51%. At the same time, 48% report mixed feelings about depending on automated systems, and open responses carry real anxiety about roles disappearing and about being asked to prove AI competence with no clear path to build it. The workforce is not resisting AI. It is navigating a transition without enough support under it.

Sixty-four percent of organizations have an active initiative to prepare their workforce for AI-driven security, up modestly from 58%. That leaves 36% without one, a figure that has barely moved even as AI has become far more embedded in daily work. Among those with programs, the common approaches are partnering with universities and online platforms (56%), mentorship (47%), and building organizational adaptability (45%). The emphasis is on access to learning more than on curriculum tied to operational need. Access and effectiveness are not the same investment.

The skills gap is not a separate workforce problem. It runs underneath every operational finding in this survey.

Conclusion

The 2026 data finds practitioners at an inflection point where the adoption decision is made, the confidence gained from deployment is real, and so are the failures, the governance deficits, and the adversarial escalation running alongside it. The question the data poses is whether organizations read all of this as a technology maturing through normal friction or as evidence that deployment has outrun readiness.

AI is producing genuine value in incident investigation, anomaly detection, red teaming, and AppSec, and practitioners are more capable with it than without. At the same time, misdirection is common, the trust gap is structurally unsolved, governance has not advanced despite a year of expanding mandates, and adversaries are keeping pace.

The next 12 months will test whether organizations can close the readiness gap against the deployment pace they have already committed to. In practical terms, that means three investments:

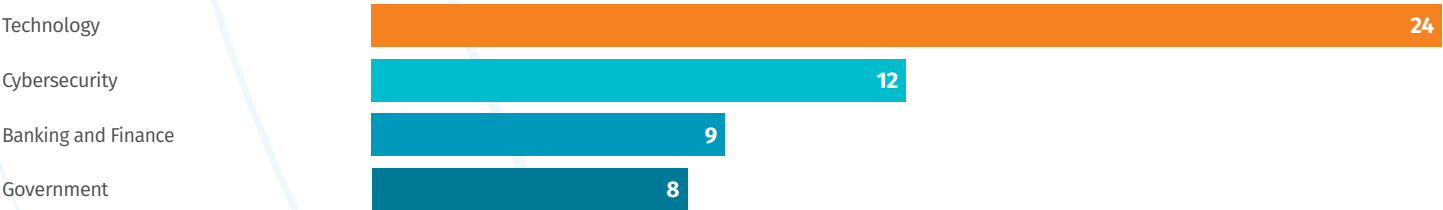
1. AI validation infrastructure, including precision, recall, and continuous comparison, rather than more deployment volume
2. Moving governance out of policy documents and into the operational controls practitioners use, with sensitive-data access and AI data exposure treated as core controls, not afterthoughts
3. Workforce development handled as an immediate operational need, not a medium-term hiring goal

Organizations that work all three at once will be in better shape than those that treated the decision to adopt as the finish line. The field has committed to AI. The harder work is making it run reliably, under governance that holds, at a pace that does not hand the adversary a standing advantage.

Survey Demographics

The 2026 SANS AI Survey received 536 responses from IT and security professionals globally. All 536 respondents confirmed their organization is currently implementing or planning to implement AI technologies Figure 9 provides the demographic details. A parallel instrument completed by 57 CISOs and senior security executives is addressed in the opening Cyber Leaders section.

Top 4 Industries Represented



Regions



Top 4 Roles Represented



Figure 9. Demographics of Survey respondents in percent

Sponsors

SANS would like to thank this survey's sponsors:





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.