

Velocity

V5 2026

Booz Allen

INSIGHTS FOR INNOVATORS

REIMAGINING
CYBER FOR A
FASTER FIGHT

IN THIS ISSUE: a16z Interview p 4 / Zero-Knowledge Proofs p 10 / Trusted Agents p 40 / Zero Trust for OT p 54

REIMAGINING CYBER FOR A FASTER FIGHT

Securing enterprises against AI threats requires disrupting operating models, enriching detection, and strengthening resilience—because attacks now unfold in minutes, not days.

By Brad Medairy and Andrew Turner

AI is now operating across the entire cyber kill chain—from vulnerability discovery and exploit development to effect delivery. This is no longer theoretical.

Consider a recent mass vulnerability exploitation against a widely used, web-based system administration tool. Threat actors identified a zero-day vulnerability in its Python script, weaponizing it to bypass two-factor authentication. This example is noteworthy as one of the first observed examples of “a threat actor using a zero-day exploit [believed to be] developed with AI,” according to a recent Google report. Evidence consistent with AI-assisted exploit development included a hallucinated CVSS score and an unusually textbook, heavily documented Python implementation—traits they said were highly characteristic of code commonly found in LLM training data.

AI-powered attacks are the new attack vector, and they are accelerating rapidly. The Five Eyes intelligence alliance recently warned that major businesses and governments were only “months” away from exposure to AI-enabled breaches. And a recent OALABS Research report found that advanced LLMs are lowering the threshold for a successful attack dramatically, finding “[i]n many cases, the attacker supplied only vague, low-skill prompts and allowed Claude to fill in the gaps: researching exposed services, identifying possible vulnerabilities, writing exploit code, validating access, and harvesting data.” Our Incident Response teams are documenting the increased pervasiveness of AI-powered attacks across both the public and private sector as well.

DIVERGING TIMELINES, EXPANDING ATTACK SURFACES

The widening gap between the pace of these attacks and the speed of response is the most urgent fault line to address in cybersecurity today. AI-enabled adversaries are demonstrating an ability to move far faster than defenders that rely on human speed, compressing the lifecycle of attacks from reconnaissance to exfiltration from weeks or days into hours and minutes.

The UK-based AI Security Institute's recent evaluation of Anthropic's Claude Mythos found that the frontier AI model was rapidly automating tasks—vulnerability, exploitation, multi-stage attacks—that would take skilled human professionals days of work. Sen. Mark Warner—citing a briefing from NSA Director Gen. Joshua Rudd—stated during a recent Senate hearing that Mythos "broke

79%

of federal cyber and IT leaders are **Extremely or Very Concerned** about adversaries using AI to accelerate cyber-attacks against their agency in the next 12-18 months.

into almost all of our classified systems, not in weeks, but in hours" during an authorized testing exercise conducted through Project Glasswing. A U.S. official separately confirmed to the Associated Press that Mythos had identified vulnerabilities in highly sensitive government systems within hours—while noting that identifying vulnerabilities and exploiting them aren't the same thing.

These advances create operational asymmetry. A lone wolf attacker can replicate the labor of multiple hackers working around the clock by leveraging AI agents, at marginal cost. Meanwhile, the defender must compete while wearing the operational equivalent of a weight vest—navigating approval chains, silos, regulations, fragmented ownership, and budget constraints.

In research conducted for Booz Allen's Vellox Striker™, which simulates AI-powered adversaries, a traditional security operations center (SOC) response spanning from initial EDR alert to incident commander approval took 45 minutes. The autonomous red team had achieved full domain compromise in six. The issue is pace, and the numbers make it clear.

But speed alone doesn't capture the full challenge. Modern enterprise environments—sprawling across cloud infrastructure, edge deployments, and increasingly, embedded AI agents—have grown dramatically more complex, and that complexity demands more instrumentation and monitoring. The drive for visibility into every layer and the expanding ability to track performance across distributed systems has produced an explosion in sensor data and telemetry, even in the absence of active threats.

AI-enabled attacks compound this further: they don't just move faster, they generate more probes, more alerts, and more simultaneous vectors. SOCs built for human-speed triage weren't designed for either the ambient volume of modern telemetry or the attack-generated noise on top of it. The result isn't heightened vigilance—it is analyst overload. Signal-to-noise ratios that were already difficult to manage are being overwhelmed, and the asymmetric, low-and-slow threats most likely to succeed are precisely the ones most likely to be missed. Defenders are buried in volume while the meaningful signal slips through.

A deficit in landscape visibility highlights an additional problem. Theoretically, AI should empower attackers and defenders proportionally—and in narrow domains it does. Microsoft used AI tools to find and patch over 500 vulnerabilities in just a few months this year, with other tech firms also reporting AI-assisted surges in patch cycles and security fixes. This suggests that defenders should have an edge: They know their own systems better than attackers do.

In reality, many organizations lack the situational awareness to make that advantage operational. Environmental knowledge is fragmented. Coverage is uneven. The intelligence needed to anticipate attacker movement—rather than simply react to it—is either unavailable or not actionable at speed. As a result, vulnerabilities linger.

These three compounding pressures—faster attacks outrunning human response timelines, SOC overload degrading detection fidelity, and insufficient landscape visibility limiting proactive defense—define the challenge security leaders must now address. The strategies that follow are organized around them directly. CISOs and other security leaders need to:

Develop Faster Responses to Attack

Build human-AI teaming models that minimize human-in-the-loop delays and enable response at AI speed.

Improve Detection Fidelity

Enrich cyber data, recalibrate signal-to-noise ratios, and rebuild SOC architectures capable of detecting more numerous and asymmetric threats.

Enhance Cyber Resilience

Adopt an "attack to defend" posture that aggressively identifies and mitigates vulnerabilities before adversaries can exploit them.

THE NEW MATH OF AI CYBER SPEED

26

SECONDS

According to an MIT/University of Naples report, an AI-supported framework discovered **13** exploit chains in **26** seconds.

34

MINUTES

In ReliaQuest's **2026** Annual Threat Report, studies of attacks found that achieving lateral movement took **34** minutes, with the fastest breakout time falling to four minutes.

1.2

HOURS

The fastest top **25%** of intrusions reached exfiltration in **1.2** hours, reduced from **4.8** hours a year ago, according to Palo Alto analysis of real-world incident response data.

10-14

DAYS

Meanwhile, the patching cadence at a typical firm might take **10** to **14** days. Organizations public and private rely on old operating systems and code bases with baked-in weaknesses, many of which may not be patchable. The programming language COBOL, for instance, is the foundation for **43%** of online banking systems and **95%** of ATM transactions—and is more than **60** years old. Layers of tech debt introduce further latency.

KICKER

The problem isn't simply attack volume. It's the shrinking decision window that legacy governance can't meet.

THE CISO'S OPERATING MODEL IMPERATIVE

AI-powered tools can accelerate detection, automate response, and surface vulnerabilities at scale—but deployed inside an outdated operating model, they will enable the wrong decisions faster. The strategies that follow require a new foundation and a new way of thinking that breaks with old assumptions and ideas underlying traditional security architecture. Technology is only part of the solution. Building this new foundation is a leadership challenge before it is a technical one.

CISOs and other security leaders must embrace this challenge. Not simply because cybersecurity is their domain, but because the changes required cut across legal, finance, compliance, IT, and business operations simultaneously. Securing agentic AI, restructuring authorization frameworks, and reorienting the workforce all require an integrator who can operate at the intersection of technical reality and organizational authority. CISOs and other security leaders arrive at this moment with exactly that mandate—and the credibility to drive the enterprise-level alignment these strategies demand. Three priorities should anchor that effort.

Reevaluate the Risk Equation

Most organizational risk models were calibrated for a different threat environment—one in which attacks unfolded over days, giving defenders time to detect, deliberate, and respond. AI has broken that assumption. When adversaries can move from initial access to full domain compromise in minutes, the window between incident and material impact may close before an organization even has visibility on the breach. SEC disclosure timelines, designed around human-speed incidents, become operationally strained. Cyber insurance policies underwritten against the same assumptions may leave organizations significantly exposed.

CISOs need to drive this recalibration explicitly, starting with an honest assessment of what the current risk model was actually built for—and where it no longer holds. That means revisiting assumptions about threat velocity, reexamining software supply chain exposure, and ensuring that zero-trust and data protection strategies are tightly aligned to crown-jewel assets. The risk model is the foundation everything else is built on. If it is out of date, every downstream decision inherits that misalignment.



Rethink Human-AI Teaming

The deeper question behind authorization and approval chains isn't simply how to move faster—it is how to rethink the relationship between human judgment and machine action in a threat environment where the adversary no longer operates at human speed. That is a more fundamental shift than process reform. It is a reconsideration of where human decision-making adds irreplaceable value, and where it has become the bottleneck that adversaries are actively exploiting.

The answer isn't to remove humans from the loop—it's to redesign where in the loop they sit. Pre-authorizing tiered response actions moves human judgment upstream, into the planning process, rather than inserting it as a gate on every containment decision in real time. Low-risk actions like quarantining suspicious files can be delegated to autonomous execution. Medium-risk actions like revoking credentials can be automated with single-human confirmation. High-risk decisions remain human—but pre-scripted against defined trigger conditions rather than improvised under pressure.

Tabletop exercises that simulate AI-speed attacks are the most effective mechanism for building this architecture: they surface gaps, clarify responsibilities, and build the cross-functional trust that allows leadership to delegate authority before a crisis demands it.

This is ultimately a workforce question as much as a policy one. Analysts who once monitored and triaged will increasingly serve as orchestrators and decision authorities—overseeing AI agents, validating outputs, and exercising judgment on the calls that machines shouldn't

make alone. Building that capability requires deliberate training, cultural change, and leadership investment. The organizations that treat Human-AI teaming as a core competency—rather than a tool adoption—will be the ones that close the speed gap.

Communicate the Resourcing Reality to Leadership

The changing threat landscape has resourcing implications that boards, political appointees, and senior leadership need to understand explicitly — and CISOs need to be the ones making that case. This isn't a technology refresh cycle. It is a structural shift in the scale and velocity of the challenge.

Vulnerability management alone illustrates the point. In just its first month, Anthropic's Project Glasswing uncovered roughly 10,000 critical and high-severity new vulnerabilities across its own and partner efforts. Addressing that backlog while simultaneously defending against active threats requires sustained investment over the next two to three years—not a one-time budget line. Cyber insurance coverage, board-level risk appetite, and compliance frameworks all need to be revisited against current threat assumptions, not the ones embedded in last year's planning cycle.

CISOs who frame this conversation in business or mission terms—exposure to material breach, insurance gaps, supply chain liability, the compounding cost of reactive versus proactive posture—will be better positioned to secure the operating model changes and resources the strategies that follow require.

WHAT FEDERAL CYBER LEADERS ARE TELLING US

Booz Allen surveyed more than **100** federal IT and cybersecurity decision-makers—spanning civilian agencies, DoW branches, and the intelligence community—about their fears and readiness for AI-driven cyberattacks. Most respondents directly shape or contribute to cyber and AI strategy at their agencies. What they told us should concern every security leader, in government and commercial enterprise alike.

THE CONCERN ISN'T HYPOTHETICAL

Nearly four in five federal cyber leaders—**79%**—say they are extremely or very concerned about adversaries using AI to accelerate attacks in the next **12 to 18** months. And the worry isn't abstract. Across five specific AI-enabled threat categories, concern levels clustered tightly above 80%: adversarial AI that manipulates or corrupts agency systems led at **85%**, followed closely by AI-accelerated vulnerability exploitation (**84%**) and autonomous attack agents (**83%**).

When asked to name the single threat capability that worries them most, leaders pointed to AI-accelerated vulnerability exploitation—the shrinking window between discovery and weaponization—as their sharpest concern.

That finding maps directly to the speed asymmetry at the center of this problem. The pace of AI-enabled attacks is what keeps federal cyber leaders up at night.

THEY KNOW WHERE AI CAN HELP—AND AREN'T SURE IT WILL BE ENOUGH

Federal leaders see meaningful defensive upside from AI, particularly in threat detection and behavioral analysis, which nearly a third identified as the area of greatest near-term benefit. Automated incident response and vulnerability management followed, tied at **20%** each. The prioritization is telling: leaders are targeting AI precisely where it can compress time—detection, response, exposure management—which is exactly where human-speed operations fall shortest.

Yet optimism is tempered by doubt. Asked whether AI-powered defenses can keep pace with AI-enabled attacks, only **36%** said yes. Thirteen percent said no. The most striking figure: **50%** are simply unsure—a level of institutional uncertainty that itself signals risk.

READINESS LAGS THE THREAT

The readiness picture is sobering. Just **6%** of respondents say their agencies are fully prepared, with AI-based cyber defenses actively deployed. A quarter of respondents report being substantially prepared, with capabilities underway. Nearly half—**48%**—are only partially prepared, and one in five is still in early-stage assessment.

The barriers driving that gap deserve close attention. The leading obstacles aren't budget or technology in isolation—they're structural. Integration with existing security infrastructure tops the list at **56%**, followed by concerns about tool reliability and accuracy at **52%**. Lack of skilled personnel (**36%**) and budget constraints (**33%**) trail behind. Policy and regulatory uncertainty, cited by **17%**, may be understated given how directly compliance frameworks complicate autonomous response decisions.

Read together, these barriers describe an organization that can see the threat clearly, has begun investing in the tools, and is being slowed—above all—by integration complexity and the absence of clear decision frameworks for when and how AI systems are authorized to act.

How concerned are you with each of the following AI-enabled threats? % very or somewhat concerned

Adversarial AI—attacks that manipulate, blind, or corrupt AI systems	85%
AI-accelerated vulnerability exploitation—near-zero window to respond	84%
Autonomous cyberattack agents—initiate, adapt, persist without intervention	83%
AI-generated influence and deception—phishing, deepfakes, synthetic identity	81%
AI-enabled insider threat amplification—nation-state-level damage threshold lowered	81%

WHAT THIS TELLS THE WHOLE MARKET

Federal cyber leaders are leading indicators. Commercial security leaders across critical infrastructure, financial services, and manufacturing report the same patterns in our engagements: high alarm, uneven readiness, governance and integration as the primary friction points. The battlespace—and bottlenecks— are shared.

There are unmistakable signs of progress: a quarter of federal agencies already have substantial AI defense programs underway, and leaders are directing investment toward exactly the right capabilities. But the dominant finding is a **50**-percentage-point gap between concern and confidence. There are barriers that no single agency, and no single enterprise, can close alone.

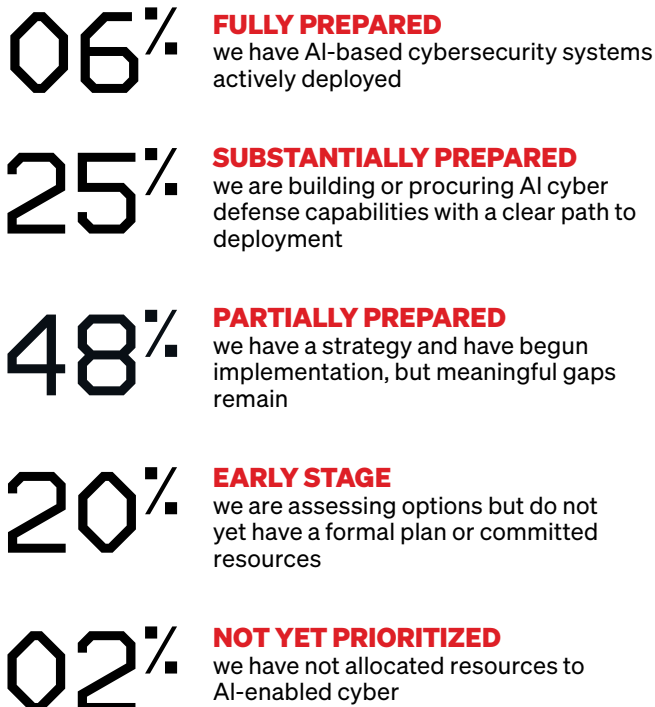
A NEW CYBER OPERATIONS MODEL FOR AI SPEED

The challenges are clear. Attacks are outrunning human response timelines. SOC architectures are drowning in telemetry while meaningful signals slip through. And defenders lack the landscape visibility to get ahead of threats rather than simply react to them. What follows is a framework for addressing all three—not as a future-state aspiration, but as an operational imperative that leading organizations are building now.

The vocabulary we use to describe offensive and defensive tactics remains in play. Adversaries employ reconnaissance, vulnerability discovery, and exploit development to attack. Defenders rely on identity controls, segmentation, and zero-trust architecture to respond.

But AI has so warped the speed and scale at which this exchange plays out that a reimagination of foundational capabilities is called for. The three strategies below rewire operating principles to account for the new attack velocity, the analyst overload created by an eruption of data and telemetry, and a new era of managing round-the-clock risk.

How would you characterize your agency's readiness to employ AI-powered cyber defenses?



Develop Faster Responses to Attack

Closing the speed gap requires moving human judgment upstream and machine execution downstream. Autonomous agents can now increasingly handle endpoint detection, triage, and initial response—from alert prioritization to recommended containment action—surfacing only the decisions that genuinely require human authority. For example, testing by Booz Allen shows that automated malware analysis can compress work that would take human teams up to 10 days into just minutes, illustrating how AI-driven triage can reclaim significant time and accelerate response. Smart data pipelines that enrich telemetry with threat intelligence accelerate that further, giving agents the context to act rather than simply flag.

The authorization framework built in the operating model phase is what makes this executable. Pre-authorized, tiered response thresholds—developed across legal, compliance, and leadership before any crisis requires them—mean that when an agent recommends containment, the decision has already been made. The goal is to move human decision making into the planning cycle where it strengthens response, rather than leaving it as a gate on every containment action in real time.

Cross-sector coordination multiplies the effect. Shared threat indicators allow organizations to synchronize autonomous responses across boundaries. Federal agencies including NSA and CISA can serve as vetted intelligence sources that increase both the speed and confidence of autonomous action.

Improve Detection Fidelity

Speed of response is only valuable if detection is accurate—and right now, for most organizations, it isn't reliable enough. The sensor-overload problem identified earlier isn't simply a volume issue. It is an architecture issue. SOC's built to funnel telemetry to human analysts in sequence weren't designed for the ambient data modern enterprises generate, let alone the additional volume AI-enabled attacks produce on top of it.

Rebuilding detection fidelity requires two simultaneous moves. The first is data enrichment: smart pipelines that contextualize raw telemetry against threat intelligence, asset criticality, and behavioral baselines before it reaches an analyst. This is what converts volume into signal. The second is structural: fusing SOC and NOC functions to eliminate the handoff delays that create exploitable windows and deploying AI agents to handle routine alert clearance so analysts can focus on the detections that require genuine judgment.

Zero trust architecture is the structural complement to both. Continuous verification of devices and identities, least-privilege access, and granular segmentation constrain what an attacker can reach even after a breach. As AI agents themselves become an expanding attack surface, zero-trust principles apply to machine identities with the same force they apply to human ones.

Enhance Cyber Resilience

The most durable advantage in AI-speed cybersecurity isn't faster reaction—it's making reaction less necessary. An "attack to defend" posture is the operational expression of this principle, and it is increasingly recognized as a foundational competency rather than an advanced capability. Frameworks like Continuous Threat Exposure Management (CTEM) formalize the approach: rather than treating vulnerability management as a periodic exercise, organizations continuously assess their own attack surface from the adversary's perspective, prioritizing exposure by actual exploitability and business impact rather than raw severity scores. AI accelerates every stage of that cycle—discovery, prioritization, validation, and remediation—at a pace and scale no human team could sustain independently.

Enhanced penetration testing is where this becomes most concrete. Agentic red-teaming programs, in which AI-powered tools continuously probe for weaknesses across the environment, are generating the kind of systematic, attacker-perspective visibility that traditional annual pen tests could never provide. Leading cybersecurity programs are already orchestrating dozens to over a hundred AI agents for this purpose—not as a one-time audit, but as an always-on capability that maps exposure before adversaries can exploit it. As AI compresses the window between vulnerability discovery and weaponization, the ability to find weaknesses first is no longer a best practice. It is the primary competition.

This logic extends upstream for organizations that develop and ship software. Human-AI collaboration in secure code development reduces vulnerability introduction before production rather than remediating it after deployment—a significantly more cost-effective posture that treats resilience as a design principle rather than a remediation task.

The workforce dimension can't be separated from the technical one. As vulnerability discovery and triage become increasingly automated, analysts shift from monitors to orchestrators—overseeing agent outputs, validating findings, and exercising judgment on the decisions machines shouldn't make alone. Like the Air Force's Loyal Wingman program, in which semi-autonomous systems fly alongside human operators, AI agents are force multipliers. Building the human capability to direct them effectively—to interpret what the agents surface, challenge what they miss, and act decisively on what they find—is as important as the agents themselves.

SPEED READ

AI-powered cyberattacks have compressed the cyber kill chain to minutes, overwhelming human-paced decision structures, drowning SOC's in noise, and exposing visibility gaps that let adversaries move faster than defenders can respond.

Operating at AI speed means disrupting the operating model: shifting human judgment upstream, empowering AI speed containment, enriching telemetry to restore signal to noise, and adopting continuous "attack to defend."

Exclusive survey of more than 100 federal cyber leaders quantifies the urgency: 79% are extremely or very concerned about AI-enabled attacks in the next 18 months, yet only 6% consider themselves fully prepared—a gap requiring immediate action.

CONCLUSION

The evidence points in one direction. AI isn't an emerging threat on the horizon—it's a disruptive reality already reshaping the cyber kill chain, compressing attack timelines from days to minutes, and widening the gap between what defenders can see and what adversaries can do. Our survey of federal cyber leaders puts a number on it: 79% are extremely or very concerned, and only 6% consider themselves fully prepared. Our work in the commercial sector finds similar sentiments. That gap between concern and readiness is the real finding—and closing it is the work.

Closing that gap requires more than deploying new tools. Risk assumptions, authorization frameworks, detection architectures, and vulnerability management practices all need to be rebuilt around the actual threat timeline—not the one last year's planning cycle assumed. The operating model has to change before the technology can deliver.

The fundamentals of cybersecurity haven't changed, and neither has the essential challenge: Find the threat before it finds you. What has changed is the speed at which that competition plays out and the organizational capacity required to win it. The frameworks exist. The strategies outlined here are being deployed by leading programs today. What determines outcomes is whether organizations treat this moment with the urgency it demands—because the adversary already is.

Brad Medairy is president of Booz Allen's National Cyber business and Andrew Turner is president for Booz Allen's Commercial Cyber business.

Survey Methodology
Market Connections and Booz Allen partnered to design an online survey of over 100 federal government decision makers/influencers, involved with IT and cybersecurity, regarding AI's impact on government cybersecurity practices. The survey consisted of 14 questions and was fielded in April of 2026. Due to rounding, responses may not add up to exactly 100%.

