



The cyber proving ground

Navigating the seams of enterprise risk and proving cyber resilience in the AI era.

- 🏆 Trusted by 30%+ of the Fortune 100
- 🏆 #1 In The Forrester Wave™ 2026
- 🏆 25M+ Exercises Run

The business isn’t waiting for cybersecurity	02
The collapse of the cyber operating model	03
Assumed capabilities vs. proven performance	04
The pitfalls of two extremes	06
The cyber proving ground: move fast, safely	07
The AI maturity curve	12
Transforming one-off exercises into an operational discipline	13
Losing control is optional	14
The cyber proving ground checklist	15
How does your organization measure up?	18

The business isn't waiting for cybersecurity

Yet, this rapid transformation has triggered a severe structural friction point. Innovation velocity is fundamentally outpacing the legacy security controls built to govern it. While the business demands the upside of AI, security leadership is left with an urgent dilemma: How do we move at pace without completely losing control? More specifically, how do we ensure AI is adopted safely, agents are built securely and with token economics in mind, and that teams are ready for AI-enabled attacks? This guide outlines the shift from basic AI adoption to comprehensive AI assurance. It establishes a blueprint for "bounded autonomy" - where speed is made safe through empirical, auditable metrics rather than static policies or assumed readiness.

Driven by the promise of unprecedented productivity, automated workflows, and rapid delivery, global enterprises are rushing to integrate artificial intelligence into the core of daily operations. From embedded developer copilots and automated triaging systems to autonomous customer service agents and executive decision-support engines, AI has crossed over from experimental pilots into live enterprise environments. Industry research indicates that 83% of enterprises are already actively utilizing AI.

■ 2026 THALES DATA REPORT

33%

of organizations have complete knowledge of where their sensitive data resides.

■ 2026 VERIZON DATA BREACH INVESTIGATIONS REPORT

67%

of users now access generative AI services from non-corporate accounts on corporate devices

■ VORLON 2026 CISO REPORT

99.4%

of security leaders experienced at least one SaaS or AI ecosystem security incident over the past year.

The collapse of the cyber operating model

For decades, enterprise security was constructed around a straightforward, foundational boundary: technology creates risk and human operators enforce the controls.

The arrival of automated enterprise AI completely shatters that assumption. Today, technology, human judgment, and autonomous AI systems seamlessly co-exist inside the exact same workflows, functioning as a singular, tightly coupled execution chain. In this modern framework, AI does not simply sit quietly beside the enterprise as a passive assistant. It actively recommends, summarizes, generates code, orchestrates operations, and increasingly mediates the high-pressure decisions that humans make. This architectural shift radically alters the enterprise attack surface across three critical vectors:

01 ACCELERATED ATTACK PATHS:

Threat actors, including state actors and criminal groups, are leveraging identical AI acceleration to lower the operational cost of reconnaissance, hyper-personalized social engineering, rapid scripting, and zero-day exploitation. The result is a harsher environment where defenders face compressed timelines and complex, synthetic threats.

02 COMPOUNDED SOFTWARE RISK:

With AI coding assistants, more people than ever can produce working code, generate scripts, and ship internal tools. This democratization dramatically expands who can introduce systemic vulnerabilities into the business, accelerating the propagation of poor quality code for AI hallucinations, unreviewed dependencies, and unsafe agent behaviors at scale.

03 THE VULNERABILITY OF THE SEAMS:

The primary threat vector is no longer found just within isolated endpoints or standalone applications. Security leaders are now forced to defend the complex, highly ambiguous seams of the organization, or the critical handoff points where humans, autonomous agents, data permissions, and machine outputs intersect.

04 EVOLVING REGULATORY AND AGENTIC FRICTION:

Managing a cross-functional automation wave across complex global business lines is an unprecedented hurdle. CISOs cannot simply act as enforcers, completely stalling business velocity. Simultaneously, passive policies are entirely obsolete against modern compliance demands. Global enterprises are now required to replace self-certification with continuous, performance-based validation to map intent and audit trails directly against a strict matrix of global AI mandates, including the EU AI Act, ISO/IEC 42001, the NIST AI Risk Management Framework, and MITRE ATLAS standards.

Assumed capabilities vs. proven performance

Many organizations are unintentionally practicing for a world that has already ceased to exist. Driven by a desire to reassure boards, regulators, and insurers, leadership teams frequently rely on static compliance checklists, annual awareness modules, and isolated tabletop exercises to claim organizational maturity.

However, macro-industry data from 2026 reveals a staggering, systemic AI readiness gap:

 THE VALIDATION AND CONTROL GAP

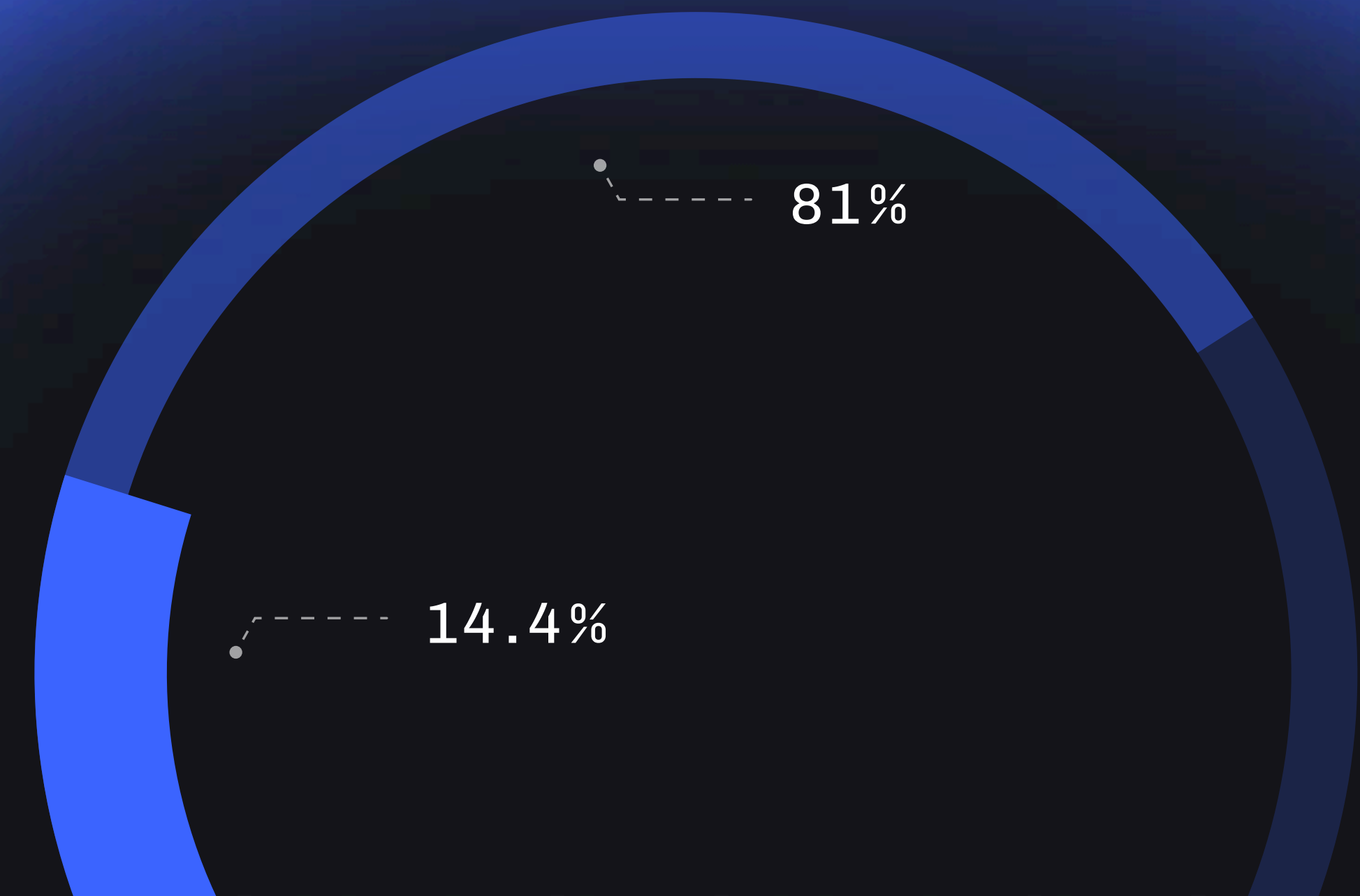
 THE DATA EXPOSURE RISK

 THE TRUE IMPACT & CRISIS COST

 THE LESSON FOR THE AGENTIC ENTERPRISE

 THE VALIDATION AND CONTROL GAP

¹ GRAVITEE STATE OF AI AGENT SECURITY 2026



While 81% of executive and practitioner teams are past the planning phase of AI deployment, a mere 14.4% have full security approval¹. This rift exists because employees are adopting AI tools faster than security teams can write policies, and autonomous AI agents are being granted real authority in production without an owner, an identity, or a documented permission boundary.



THE DATA EXPOSURE RISK

True visibility into data exposure has collapsed. Recent data reveals that only 33% organizations have complete knowledge of where their sensitive data resides². This blind spot is actively being exploited under the radar: 67% of users now access generative AI services from non-corporate accounts on corporate devices, with proprietary source code ranking as the most common data type being uploaded directly to public tools.

THE TRUE IMPACT & CRISIS COST

This underground adoption presents a live operational threat. 99.4% of security leaders experienced at least one SaaS or AI ecosystem security incident over the past year. Furthermore, organizations tolerating high shadow AI exposure are bearing the brunt of the financial fallout, facing average breach costs of \$4.63 million - roughly \$670,000 more than organizations maintaining lower exposure.

THE LESSON FOR THE AGENTIC ENTERPRISE

Checking a compliance box, completing standard training modules, or managing an unvalidated AI rollout is entirely insufficient. Policy is merely an aspiration until it is tested under realistic pressure. Training shows what people know; Immersive One proves what your people, teams, agents, and AI-enabled workflows can actually do under pressure.

² 2026 THALES DATA REPORT

The pitfalls of two extremes

Governing this shift is challenging. On one level, balancing rapid innovation with risk mitigation is a familiar friction point. For years, security has fought to shed its reputation as the organizational roadblock.

Yet, widespread automation introduces an unprecedented layer of vulnerability: as autonomous models and automated workflows are rapidly adopted, they continuously learn from human behavior. When organizations prematurely replace human oversight with automated execution to drive down overhead, the critical institutional knowledge, context, and situational judgment of the workforce are completely severed from the loop.

Without a controlled environment to validate these handoffs, global enterprises generally default to one of two extreme overcorrections. Both paths carry steep operational and financial consequences.

EXTREME 1: PRIORITIZING SPEED AND LOSING CONTROL

Organizations in this category automate prematurely before they fully map or understand their underlying workflows. They treat rapid AI rollout as progress and frame security governance as unnecessary drag.

This approach directly turns experimentation into immediate systemic risk. The enterprise suffers from unauthorized actions executed by unchecked agents, corporate data leakage via public AI platforms, and destructive workflow executions.

Many organizations attempt to reduce human overhead by prematurely cutting headcount, only to discover too late that removing humans from the loop does not remove risk from the business. The damage getting through to the business scales exponentially before controls can intervene.

EXTREME 2: PRIORITIZING CONTROL AND SACRIFICING SPEED

Terrified of the unmanaged threat landscape, these organizations resort to legacy defense tactics: banning tools, gating deployments, and demanding lengthy validation cycles.

Security cannot win by slowing the business down. When security acts as an absolute roadblock, it drives AI adoption completely underground. This forces business units to quietly deploy unsanctioned models and generative utilities - instantly creating an invisible, unmanaged shadow AI layer where proprietary code and sensitive data flow entirely outside of corporate oversight.

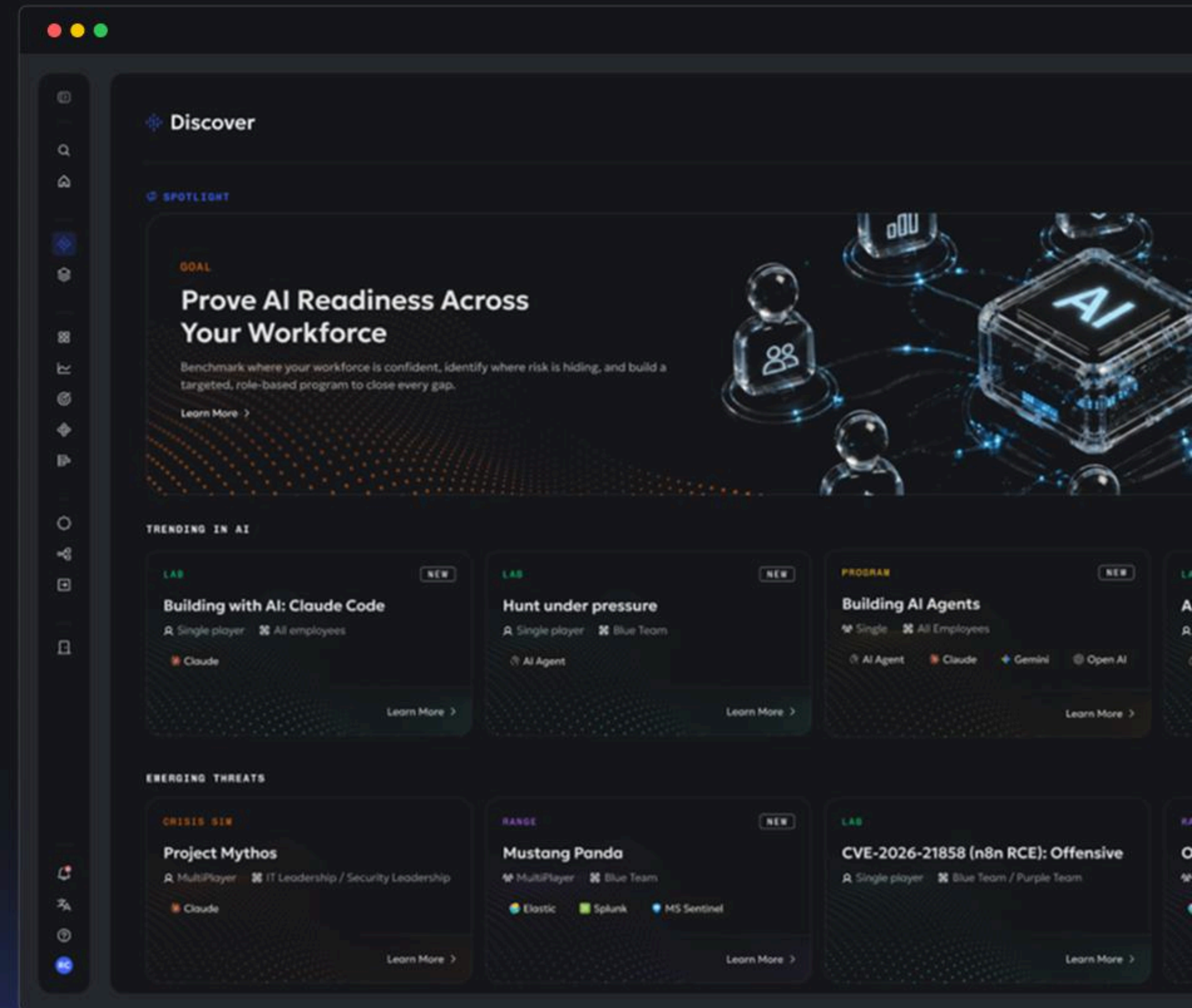
Business teams will independently bypass central security to leverage generative AI and embedded assistants to hit their productivity targets. This creates a massive shadow IT layer where sensitive data flows entirely outside of organizational visibility, maximizing the corporate attack surface while leaving security completely blind.

The cyber proving ground: move fast, safely

Security cannot protect the enterprise by acting as an organizational gatekeeper or attempting to slow the business down. To win in this new era, organizations must shift their perspective from basic AI adoption to comprehensive AI assurance.

Adoption simply asks whether the business is deploying AI tools to drive efficiency. Assurance asks whether the organization can empirically prove that its human + AI operating model can hold under pressure before it is tested for real.

This requires transforming the human role from a basic operator of tools into a rigorous governor of workflows. Humans must be trained to continuously evaluate machine outputs, challenge assumptions, manage agentic handoffs, and carry absolute structural accountability when autonomous logic breaks down. To systematically achieve this, modern cyber resilience must be anchored across four core proof pillars:



PILLAR 1:

 Adopt AI safely

Organizations must move past passive, generic policy distribution. By implementing interactive, scenario-based validation, an enterprise ensures that AI entitlement is actively earned rather than blindly assumed.

Workforce populations - from non-technical departments to developers and executives - must demonstrate practical policy compliance, data handling, and safe usage in context before being granted access to high-risk tools and models.

Pass the exercise, earn the entitlement; fail the exercise, get targeted, role-specific upskilling.

Harness the power of expert-curated guided goals with these value plans.

- All
- Cyber outcome
- Who it's for

GOAL

Build compliance and security standards capability

Turn compliance requirements into measurable security evidence. Governance teams build confidence interpreting, implementing, and demonstrating progress against the standards and frameworks that...

3 activities  Governance

VIEW DETAILS →

GOAL

Build offensive AI security capability

Strengthen your ability to test and challenge AI systems before attackers do. Teams practice AI red teaming, prompt injection, LLM risk discovery, and guardrail validation through realistic offensive scenarios.

2 activities  Offensive

VIEW DETAILS →

GOAL

Build full-spectrum

Develop end-to-end pe infrastructure, and ent reconnaissance, exploi

2 activities  Offens

VIEW DETAILS →

GOAL

Improve resilience against relevant threat actors

Prepare teams to defend against the adversaries most relevant to your organization. Practice real-world tactics, techniques, and attack chains to improve detection, response, and operational readiness.

4 activities  Defensive

VIEW DETAILS →

GOAL

Improve workforce cyber safety

Reduce everyday cyber risk by improving how employees recognize and respond to common threats. This goal builds safer behaviors around phishing, social engineering, authentication, data handling, and secure...

3 activities  Everyone

VIEW DETAILS →

GOAL

Improve AWS de

Improve cloud detecti Security teams practic cloud events, and resp

3 activities  Defens

VIEW DETAILS →

GOAL

Improve endpoint and DFIR capability

Strengthen investigation and forensic response across endpoint, network, malware, memory, and log evidence. DFIR teams practice realistic scenarios to improve speed, confidence, and evidence quality...

1 activity  Defensive

VIEW DETAILS →

GOAL

Improve organizational SIEM proficiency

Increase the value your organization gets from its SIEM. Analysts and detection engineers improve alert triage, query writing, and detection quality so incidents can be identified and escalated with greater...

3 activities  Defensive

VIEW DETAILS →

 Defend at AI speed

Technical teams must be able to keep pace without blindly surrendering critical judgement to an unverified algorithmic black box. Engineers need to catch AI-generated vulnerabilities, insecure dependencies and unsafe software patterns before they reach production. Security teams need to detect, decide, escalate, and respond under AI-speed pressure.

Utilizing advanced Live-Fire Cyber Ranges, Dynamic Threat Range scenarios, and hands-on technical exercises, organizations can stress-test builders and defenders against real-world adversary behavior. This generates clear, empirical telemetry mapping secure development capability, detection speed, escalation quality, and response accuracy against operational baselines.

Dynamic Threat Range



Exercise Your Blue Team

Empower your teams to defend real infrastructure against sophisticated live-fire attacks using your preferred SIEM.



Agentic

Exercise Your Red Team

Evaluate your teams' ability to leverage AI-powered pentesting tools against realistic target environments.



Exercise Your Developer

Stress-test your teams' ability to wield the software development lifecycle.

Your organization's exercises



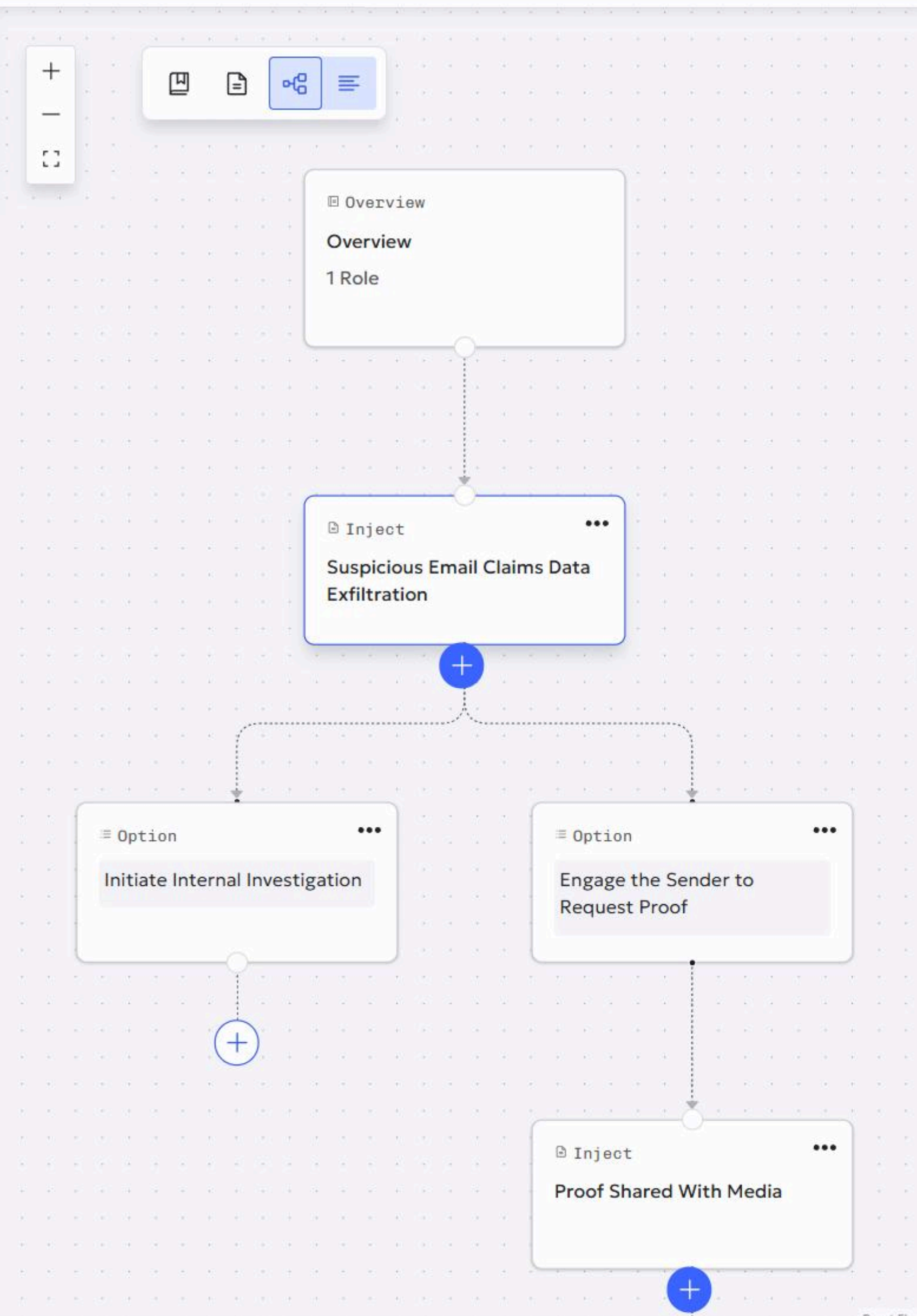
PILLAR 3:

Validate AI agents

AI agents are moving from controlled experiments into operational workflows. They're being used to triage alerts, recommend fixes, investigate incidents, test defences, and accelerate decisions across SOC and SDLC environments. But an agent used in a demo isn't the same as an agent that can be trusted under pressure.

Organizations need to validate where agents perform, fail, when humans need to intervene, where guardrails hold. As these systems scale, leaders must ensure that operational value justifies the cost. Agentic workflows must be tested for accuracy, safety, governance, auditability, human override, and token-spend efficiency before they're trusted in live environments.

Using agentic evaluation environments, teams can test autonomous behavior against realistic scenarios without introducing risk. They can compare agent performance against human baselines, expose failure modes, and identify the conditions under which agentic workflows are ready to scale. The promise of AI agents is speed. The requirement is proof.



Inject

Details

Options

Feed items

Title

Suspicious Email Claims Data Exfiltration

Inject date (Optional)

Thu, 19th Feb 2026

Time of inject (Optional)

05:00

Description

677 / 6000 characters

T ▾ B I </> ≡ ▾ ≡ 99    

What's happened

The IT service desk has received an unsolicited message from an actor identifying themselves as **BlackVeil**. The group claims to have gained unauthorised access to one of the organisation's production databases and exfiltrated a sample of customer records.

The demand

- **Amount:** 2.5 million USD, payable in Monero
- **Deadline:** 48 hours
- **Threat:** Data published to BlackVeil's leak site if unpaid

What we know, and don't

- The only "evidence" provided is a claim of *10GB of customer data*. No sample or proof has been supplied.
- The IT team has **not** yet validated the claim, attributed the actor, or notified leadership.

The IT team will soon forward the email to you. Check your

PILLAR 4:

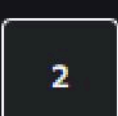


Lead through crisis

When an AI-enabled workflow experiences an unmanaged failure, an over-permissioned agent executes an unauthorized action, or a model output triggers a severe regulatory violation, the board does not want to know if your team completed a video course.

They demand auditable, courtroom-ready evidence that the business can operate, decide, disclose, and recover under pressure.

Scalable simulations push executive teams through current, hyper-realistic crisis paths, generating automated After-Action Reports that transform raw operational execution into definitive corporate governance.



Systems

Tooling

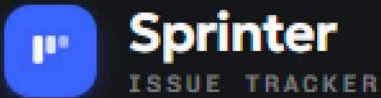
- Gitea Open
- Blossom App Open
- Coder Open
- Sprinter Open

Show briefing

Gitea

Blossom App

Sprinter Connected



FILTER All assignees All statuses All priorities Default order

Security Tickets

TO DO 6 IN PROGRESS 2 BLOCKED 0

BUG-1510

Broken Access Control on Feedback

Any authenticated caller can read another user's feedback by supplyin...

Bug Unassigned

BUG-1511

Spoofed Feedback Sender

The feedback submission route trusts a sender_id from the request body...

Bug Unassigned

BUG-1512

Excessive Time-Off Allocation

The time-off request route accepts a request whose start date is after its...

Bug Unassigned

BUG-1515

Excessive Data Exposure on Login

The login route returns the auth service's user record verbatim, leakin...

Bug Unassigned

BUG-1517

Server-Side Request Forgery in Documents

The document list route concatenates an unvalidated user_id into an intern...

Bug Unassigned

BUG-1518

The AI maturity curve: from tactical labs to continuous assurance

A common roadblock for organizations accelerating their AI journey is implementation anxiety. Prospects worry that building an empirical defense model is too complex, while existing customers frequently get stuck using platforms tactically - focusing on isolated product usage rather than system-wide strategy. True resilience is an evolutionary path. Organizations must scale their maturity over time across three distinct phases:

PHASE 1: BASELINE FLUENCY

Resilience begins by validating individual capability. Utilizing role-specific labs and targeted technical tracks, organizations establish an initial performance benchmark across distinct workforce populations and individual technical teams. Rather than relying on passive learning, this phase grounds the workforce in the core competencies required to interact with AI securely relative to their specific business functions.

PHASE 2: SYSTEM-WIDE RESILIENCE

Individual skills must eventually be tested as a unified defense network. In this phase, organizations graduate from siloed exercises into repeatable, cross-functional cyber drills and live-fire crisis simulations. This layer measures the system as a whole, tracking real-world response speed, decision-making accuracy, escalation quality, and team coordination under intense operational pressure.

PHASE 3: AUTOMATED AI GOVERNANCE

At peak maturity, validation becomes continuous, objective, and executive-ready. The enterprise leverages advanced capabilities like the AI Agentic Harness to stress-test automated security workflows against live adversary behavior, utilizes Engineering Sandboxes to validate developer context, and tracks framework-aligned peer benchmarking. This phase transforms raw performance data into defensible, board-ready evidence of true cyber readiness.

Security leaders do not succeed by buying one more detection platform or writing another theoretical text policy. The mandate is to build a disciplined, repeatable operating model that makes inevitable machine and human mistakes smaller, rarer, faster to detect, and cheaper to survive. Moving along this trajectory requires moving away from passive video libraries and single-module usage into a repeatable cyber resilience program. Mature organizations must combine individual technical capabilities with system-wide simulations to generate continuous, board-ready evidence.

Transforming one-off exercises into an operational discipline

The greatest risk to any security validation program is treating readiness as a static, one-off project rather than a continuous discipline. To maximize the value of a cyber proving ground, organizations must embed rigorous testing directly into their ongoing risk management lifecycle.

Immersive One transforms this vision into a scalable reality, serving as the continuous proof layer that unifies people, workflows, and automated agents. This continuous assurance model is sustained through three core capabilities powered by the platform:

SELF-SERVICE PROGRAM DESIGN

The risk landscape changes weekly, and security programs must adapt instantly. Utilizing a conversational AI Program Builder, leaders can describe a specific strategic goal - such as adjusting to an emerging vendor AI rollout or strengthening a team's proficiency on a new tool - and the platform will instantly curate, select, and sequence a tailored security program in seconds.

ADAPTIVE THREAT ALIGNMENT

Practicing against yesterday's curriculum leaves an enterprise entirely exposed to today's threats. The proving ground updates automatically as threat intelligence evolves, ensuring that technical teams are actively rehearsing against current adversary behaviors - such as Scattered Spider techniques - rather than static, outdated CVE repositories.

CONTINUOUS COMPLIANCE MAPPING

Audit readiness should never be a frantic scramble preceding a regulatory deadline. The platform automatically maps ongoing workforce and technical performance metrics directly to global cybersecurity standards like NIST CSF, ISO 27001, DORA, and OWASP, alongside the critical AI governance pillars of ISO/IEC 42001, the EU AI Act, NIST AI RMF, and MITRE ATLAS.

For highly regulated global financial institutions, the platform dynamically aligns live performance data with the operational resilience and governance baselines established by the Cyber Risk Institute (CRI). This gives compliance officers, insurers, and regulators an uninterrupted pipeline of definitive proof that your controls, people, and automated workflows are consistently functioning under pressure.



Losing control is optional

For modern enterprise leaders, artificial intelligence adoption is inevitable. The strategic choice, however, is establishing how they will govern and demonstrate control of the resulting risk.

Organizations that rush to over-automate, under-govern, cut human oversight prematurely, and treat rollout as progress will learn expensively that removing humans from the loop is not the same as removing risk from the business. The resilient enterprise embraces bounded autonomy. This means establishing explicit machine boundaries, mandating highly observable workflows, enforcing earned access, and maintaining absolute human authority at the moments that matter most.

By establishing a dedicated cyber proving ground across people, workflows, agents, and leadership decisions, security leaders stop guessing at capability and start generating empirical proof. We make speed safe, reduce the blast radius of inevitable operational errors, and provide the defensible evidence leadership needs to confidently accelerate into the AI enterprise.

The business isn't waiting for cybersecurity, and assumed readiness is no longer enough to satisfy your board, regulators, or insurers.

With Immersive One, you can move away from passive compliance checklists and turn raw operational performance into definitive evidence. Make speed safe, protect the seams of your risk environment, and accelerate your AI transformation with absolute confidence.

The cyber proving ground checklist

Assessing your AI assurance and readiness

To successfully balance rapid innovation with absolute security control, organizations must move away from assumed readiness and graduate to empirical proof.

This checklist serves as an operational diagnostic for both prospects and existing customers. Use it to audit your current capabilities, identify structural blind spots across your human + AI operating model, and map out your next steps toward continuous assurance with Immersive One.

01 WORKFORCE POSTURE & ENTITLEMENT

Eradicate the massive exposure of shadow AI adoption by replacing static acceptable-use text policies with role-specific, interactive validation. Ensure your non-technical staff and business teams use enterprise copilots safely without leaking sensitive corporate data into public models.

- ☐ We have clear, real-time visibility into which employees are utilizing specific public and embedded generative AI tools.
- ☐ Our acceptable-use policy is actively verified through role-specific, scenario-based testing rather than assumed via passive compliance video modules.

THE IMMERSIVE ONE EDGE :

Earned entitlement: Access to high-risk tools and models is tied directly to an AI Knowledge Pass, meaning users must pass an interactive exercise to unlock role-specific AI access.

12 SECURE SOFTWARE ENGINEERING

Prevent developer assistants from introducing structural risk into production code. Enable software builders to maximize the productivity of generative coding while systematically catching brittle assumptions and unreviewed dependencies at scale.

- ☐ We actively test our software engineers on secure AI development, including retrieval pipelines, model context, and guardrail validation.
- ☐ We can measure whether developers using coding copilots are inadvertently introducing insecure logic or unreviewed third-party packages into our repository.
- ☐ We pressure-test in engineering sandboxes: Software builders enter isolated, live code testing environments to actively stress-test and validate custom security guardrails before deployment.

THE IMMERSIVE ONE EDGE:

Validate your autonomous security systems inside the AI Agentic Harness before granting them production authority. Immersive One replaces blind trust with empirical evidence, enabling technical teams to map failure modes, establish human override boundaries, and audit the token economics of agent deployments before they go live.

13 AGENTIC SOC & OPERATIONAL DEFENSE

Ensure technical teams and autonomous security agents can detect, investigate, and mitigate compressed, machine-speed attack paths without surrendering critical defender judgment to an unverified algorithmic black box.

- ☐ Our SOC analysts actively practice triaging and responding to compressed, AI-powered threats, and synthetic social engineering.
- ☐ We have benchmarked our defensive workflows to understand how the handoff points between human operators and automated security tools perform under pressure.
- ☐ We conduct regular agentic evaluations: Custom security agents are exposed to live adversary behavior to empirically track execution success, logic boundaries, and token-spend efficiency.

THE IMMERSIVE ONE EDGE:

Validate critical human-in-the-loop capabilities across your SOC and engineering teams, proving your defenders and developers possess the precise interrogation skills and contextual judgment required to verify autonomous AI decisions before they impact production.

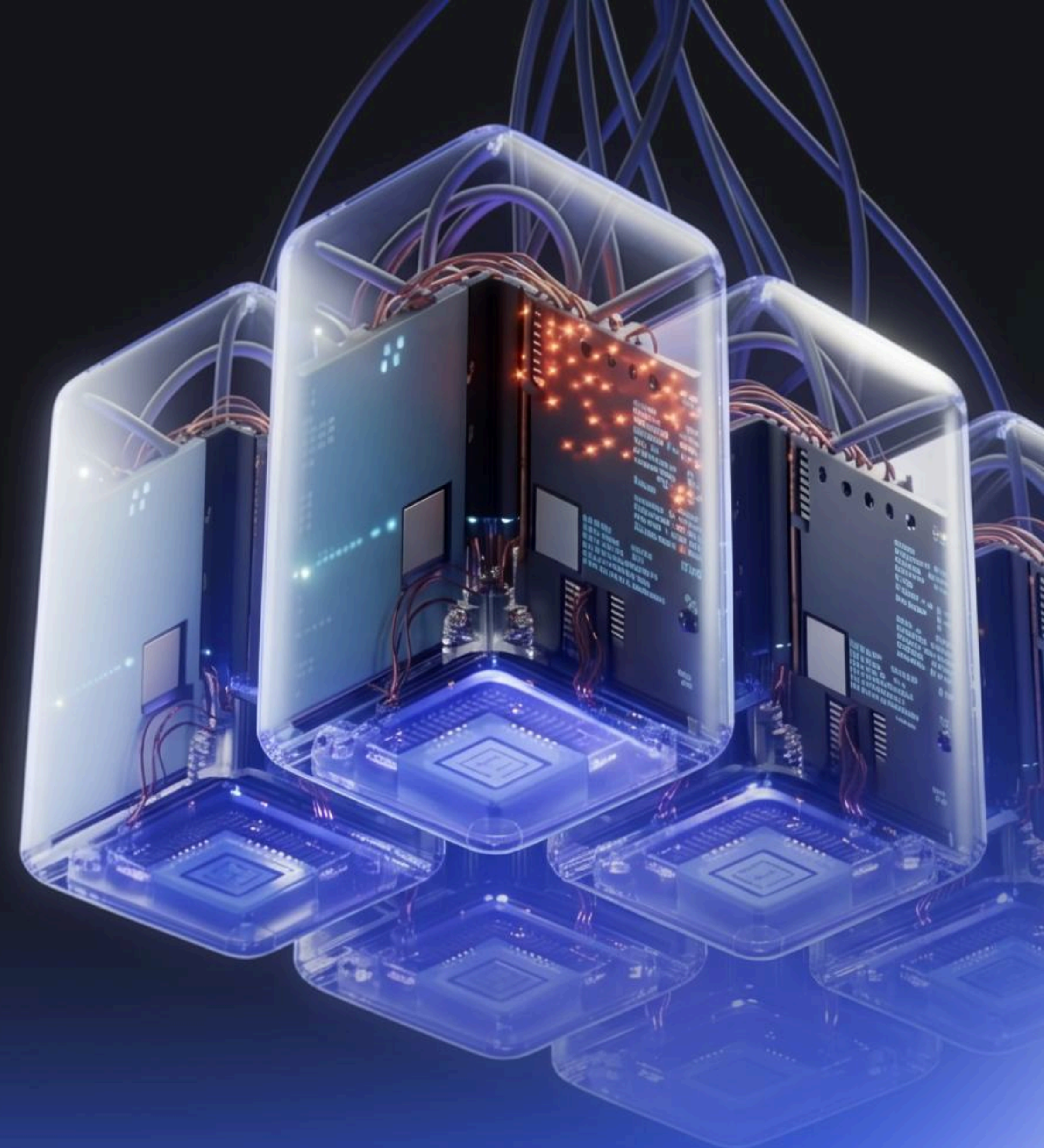
4 LEADERSHIP DECISION-MAKING & CORPORATE GOVERNANCE

Rehearse high-stakes, boardroom-ready decision-making during an AI-accelerated operational failure, and generate the exact empirical performance metrics required to satisfy regulators, insurers, and the board.

- ❑ Our executive and leadership teams have rehearsed the legal, operational, and regulatory disclosure decisions required when an AI workflow fails or creates corporate risk.
- ❑ We can present our board and insurers with defensible, auditable evidence of actual performance under pressure, rather than static maturity scores.
- ❑ We run regular leadership simulations: High-fidelity tabletop exercises are scaled across the enterprise, generating automated After-Action Reports (AARs) to map skill gaps, decision accuracy, and operational response speed.

THE IMMERSIVE ONE EDGE:

Deploy targeted crisis simulations tailored specifically to your organization's unique AI risk profile. When every employee is an AI-enabled builder, Immersive One ensures your leadership and technical teams are operationally prepared to neutralize complex, fast-moving threats.



How does your organization measure up?

In our experience across global enterprise environments, most organizations discover they are missing a significant portion of these Proving Ground capabilities. Practicing for a world that has already changed leaves critical vulnerabilities at the highly exposed seams where your staff, code, and vendors intersect. If you checked fewer than three boxes or rely on manual, one-off exercises, your organization is at risk of sacrificing speed for control or vice versa. In reality, you don't need to choose one over the other.

Don't wait for an operational failure or a regulatory audit to discover where your control boundaries break.

[See the platform in action >](#)

•••

Welcome back

Early Access

Get started with goals

Turn your strategy into measurable goals

Combine strategic and technical training in one view.

Track the unique performance data your team needs.

Pick an expert-led template or build your own

Browse goals

RECOMMENDED GOALS

Build enterprise-wide AI literacy and governance capability

Build a workforce that can use AI confidently, responsibly, and securely. Employees gain the judgment to recognize where AI adds value, where it introduces risk, and how to apply it safely in everyday work.

Build AI-first developer capa

Help engineering teams adopt AI as software delivery. Developers build o AI coding tools, designing safer LLM applying responsible AI practices acr

⚡ Manager Quickstart

Managed by me

Assigned to me

0

2

Assignments

Assign collections of labs or exercises to learners and teams in your organization.

Create assignment

Manage assignments

Your Learning

Try our recommendations below.

Recommended activities

Assigned to me

Challenges & Scenarios

Leaderboard

JD

John Doe

View Leaderboard >

Labs

Points

Rank

9

760

1/1

Streaks

Get started by completing a lab

Streak settings >

Recent Activity

TITLE	STATUS
Building with AI: Claude Code – Claude Skills	In Progress
AI: Demonstrate Your Skills	In Progress

Immersive News & Updates

Content Updates

New Building with AI Collection

23 Jul 2026

Platform Updates



Prove your business can move securely at AI speed.

Immersive is the cyber proving ground for the AI enterprise. We help organizations prove their people, AI agents, and leadership decisions are ready for a world where AI is used inside the business and against it.

