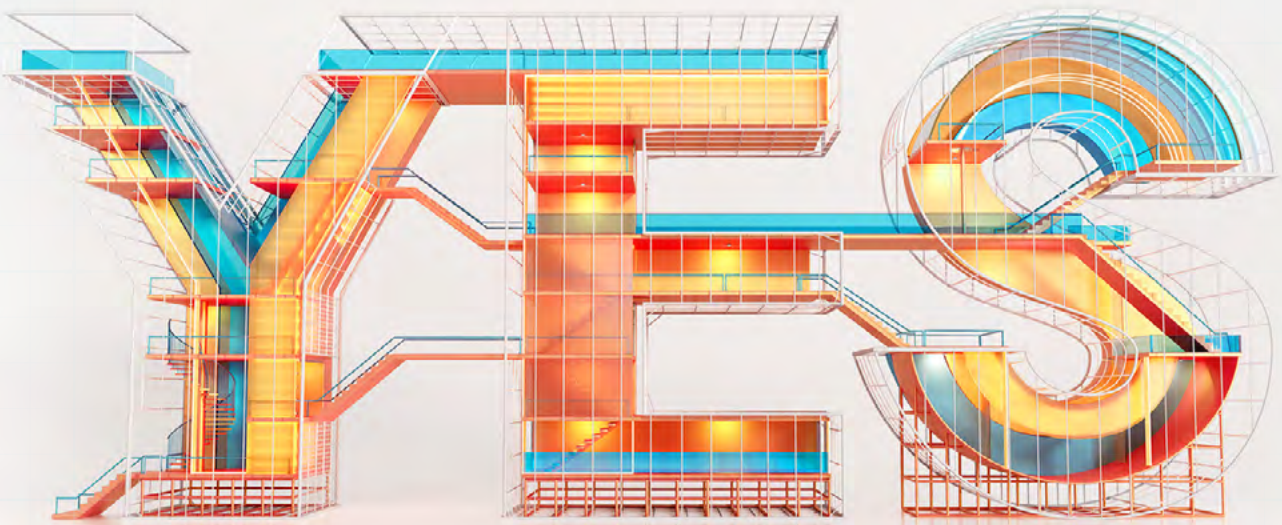




A Blueprint For the Architects of



How CISOs are Accelerating
Secure AI Adoption

Inside this report

AI is now a business imperative	4
The problem: Adoption has outrun infrastructure	5
Chapter 1. Architecting a 'Yes' for AI	6
Chapter 2. The Blueprint: Six conditions for defensible AI acceleration	10
Chapter 3. Conclusion: Build before the gap widens	17

Introduction

AI has rapidly moved from experimentation to enterprise infrastructure. Boards and CEOs are excited about its opportunity to enhance efficiency, growth and competitive advantage. How are the best Chief Information Security Officers (CISOs) in the business making sure all of this can happen safely for their organizations?

CISOs need to be able to say “yes” to their organizations’ proposed AI projects. They must do that without exposing the organization to unacceptable risk.

Netskope has just completed a listening tour, hearing from security leaders and CISOs around the world to understand what must be in place in order for their organization to adopt AI securely. This report gathers together insights gleaned from those conversations, and sets out CISO experiences in their own words, alongside their recommendations to fellow professionals.

Their evidence points to one conclusion: the next phase of AI adoption will be led by those that build the architecture to move fast in safety. With the help of their expert partners, CISOs must now become the **Architects of Yes.**

76.3

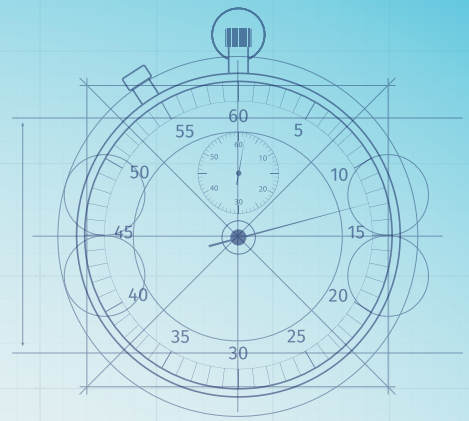
“Infosec or cyber has been built on environments and technology that are very predictable. And as we know, AI is not... From a security point of view, I think the foundations and the fundamentals of maybe 70% of what we already do is applicable. And then there’s this delta of 30%.”

– CISO, software, Australia

AI is now a business imperative

Adoption first,

readiness second



"If you look at the speed we are taking on AI, it's rather unprecedented in our business and industry. Normally we would wait for other companies to implement it, and then maybe in five or 10 years we will have a look at what worked, what didn't and so on."

– VP cybersecurity EMEA, financial services, UK

CISOs told us they have:

50-100

AI POCs in progress at once.

And POCs are running faster:

18 months > 3-6 months

Mandated from the top

72%

of CEOs say they are now the main decision-maker on AI in their organization, double the rate of 2025.

*"The impression in many companies is that **if we don't do AI, then our competitors will** and wipe us off the face of the Earth".*

– Steve Riley, Netskope's Field CTO

"The greatest pressure is everyone wants you to say 'yes, it's OK to deploy'

... We had a few cases where we pretty much were forced to say yes."

– CISO, financial services, US

Adoption has outrun infrastructure

Adoption without governance

73% of enterprises have deployed AI in 2026.*

But just **7%** of organizations are able to govern their AI in real time.

"AI is becoming embedded in enterprise workflows faster than most organizations can govern it."

– CISO, IT services, US

As one CISO put it:

"The gap is the governed execution".

Capacity and capability challenges

"The volumes that we get in terms of the issues uncovered by AI is not achievable. You can't even review whether they're valid or not with human-only personnel."

– VP cybersecurity EMEA, financial services, UK

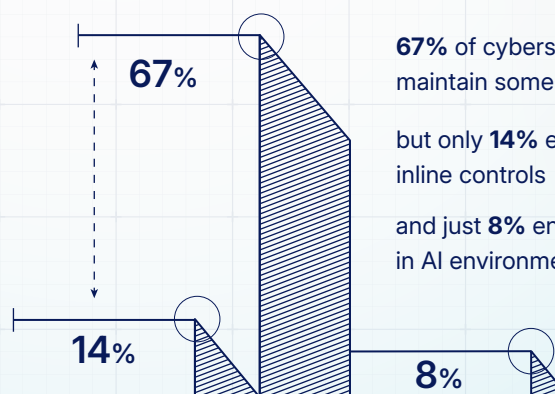
"We're not ready for the scale of vulnerabilities that are coming at us with the new frontier models."

– CISO, software, Australia

Today's governance

"It's all policy". Words on paper rather than rigorous tools.

Controls that do exist are partial and frequently bolted on after the fact.



67% of cybersecurity practitioners maintain some form of AI policy,

but only **14%** enforce through inline controls

and just **8%** enforce consistently in AI environments.*

"The cause of most breaches that relate to AI at the moment is somebody putting in a control on day one and never monitoring it again. Two weeks later the control is bypassed and the incident becomes inevitable."

– Neil Thacker, Global Privacy & Data Officer, Netskope

1. Architecting a 'Yes' for AI

The generational shift

The old narrative that CISOs are blockers to organizational innovation is years out of date. CISOs moved on to play a more positive and proactive role, even before AI arrived. But what's changed with AI is the volume and speed of requests they get. Case by case adjudication has become impossible, so CISOs are creating a better system for reviewing them.

"My role has changed from being a 'department of no' to a 'department of yes, but stay vigilant'."

– CISO, retail, UK

In this sense, the CISO's job now involves building the conditions under which the default answer should be 'yes'. Elements of their gatekeeping function still remain, especially in terms of spend and access control. But CISOs are now better thought of as an architect of secure approval, making the accelerated AI adoption curve both governable and repeatable.

"The answer is not friction. The answer is a governed path to yes."

– CISO, IT services, US

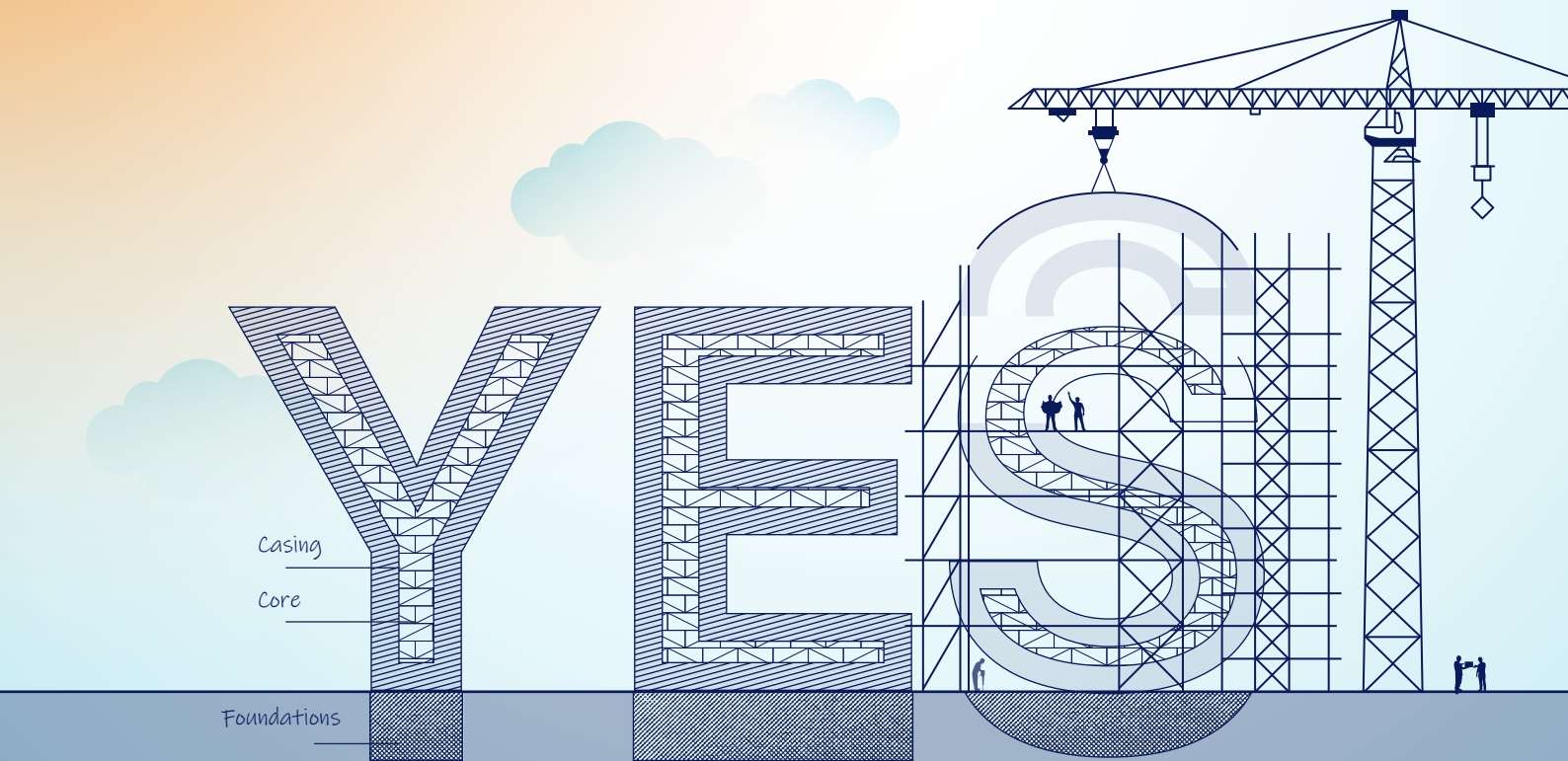
The architecture for yes

In the same way that architects must take a holistic approach to the buildings they design, CISOs tasked with securing their organizations' AI projects are working toward a consistent, in-depth process to be confident of the outcome. That was a clear finding from our interviews with security leaders at global enterprises across the US, UK, Australia, India and Israel. One CISO in the financial services sector described his approach as operating "...multiple gate checks until everyone is ready to say yes".

CISOs told us about several conditions they set for rapid yet robust AI approvals in their organizations. While not every CISO is implementing all of these approaches, together they add up to a smoother route for getting to "yes" as a default.

"If you're on an enterprise platform, new tools or applications shouldn't have to go through the security review process. That should be a 'yes' scenario."

– James Robinson, CISO, Netskope



Yes by design

Foundations

- **Clear ownership.**
Accountability for the system, its processes, and the data it handles must rest with a named individual outside the security team.
- **Business prioritization.**
The request has to have a strong business case and align to organizational objectives. Prioritization can be determined by business leadership, or via the security team's own scoring method, but either way it's essential in order to filter hundreds of AI-related requests down to a handful of top priorities.

Core

- **Assessed architecture.**
There must be clarity on how the AI system is built and what it connects to, including the models, APIs and any agents operating behind it.
- **Controlled inputs and outputs.**
Both ends of the interaction need enforcement. The prompts and uploads going in, and responses coming out (where hallucination is the key risk).
- **Green-light systems.**
By operating only a small set of approved enterprise platforms, security teams can let anything inside them proceed on acknowledgment alone so their review capacity is spent only on what's genuinely new.

Casing

- **Evidence retained.**
The decision, the risk being accepted, and who accepted it all must be documented in writing so that the process can be audited later.

The limits of the CISO's role

Risk acceptance is the business's responsibility, not the CISO's individual responsibility. In their role as an architect, the CISO structures the process to make the decision explicit to all parties, document its potential implications, and ensure the risk owner is aware of what they are agreeing to.

CISOs are not giving permission for a particular tool to be used by the business. They are instead creating the conditions for its risks to be understood and weighed, and for the relevant business stakeholder to make an informed judgment about it. CISOs' input helps influence that thinking, but ultimately it's always a corporate decision.

One CISO gave the example of a leader who wanted a new application added to his team's tech stack: "I went back to the leader and said, 'Will you accept the risk then, if this organization gets breached, compromised or goes out of business'? His response was immediately: 'Well, how do we get it through so that it's an acceptable risk?' That changed the conversation."

CISOs are not giving permission for a particular tool to be used by the business. They are instead creating the conditions for its risks to be understood and weighed, and for the relevant business stakeholder to make an informed judgment about it. CISOs' input helps influence that thinking, but ultimately it's always a corporate decision.

Telling the right story

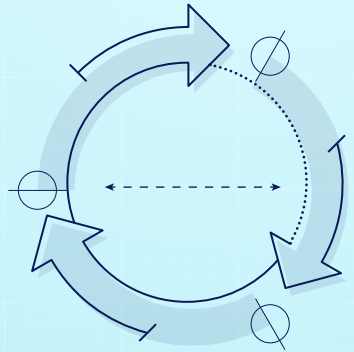
Business fluency is a key aspect of structuring this process successfully. Framing outcomes and risk through a business lens, instead of an IT or security lens, is "how you win hearts and minds" across the C-suite, as one UK-based CISO explained.

But while capable security leaders will understand the business context, many still find it difficult to explain risk in board-friendly language. "I think the CISOs that are going to have a hard time are the ones who struggle to articulate the risk and communicate it," one leader said.

A former practitioner turned CISO advisor recommended that security professionals work on their storytelling skills. He suggested talking about risk reduction, faster approvals, and measurable proof about what's blocked, what is accepted and what value is enabled or curated.

A yes that has to hold

An architect is very different from an approver.
The latter looks at each new request in turn.
The former creates a **repeatable process** that works
over a concerted period of time.



Building for consistent approval can be challenging at a time when AI services change weekly. Architecting to be able to say yes creates a framework for a standing commitment from the security team, not a single decision, therefore controls must evolve with the service they're governing. One CISO observed that "You almost need a guardian around each AI service to monitor it and adjust based on its activities".

As another CISO put it, this makes governance a "dynamic, cyclical, adaptive process," owned across the business rather than attributed to security alone. The goal is clear: "Every AI approval should be observable by design, governed at scale and enforceable before execution."

"When the mechanisms are in place, have been practiced well and feel like they become routine, then it's going to be easier for CISOs to say yes to new projects and initiatives."

– Steve Riley, Field CTO, Netskope

The previous chapters set out CISOs' experiences of AI adoption, the problems it causes for their infrastructure and controls, and the way they are shifting their role from traditional gatekeeper to an architect of governed and defensible AI acceleration. The next chapter contains CISO recommendations for putting what they've learned into practice.

2. The Blueprint: Six conditions for defensible AI acceleration

The maturity of AI security approaches varies drastically across organizations, and in most of them is still a work in progress.

Here we bring together lessons from enterprise CISOs around the world. We outline the six key conditions and suggested actions they proposed for making AI projects work fast, but also in a secure manner.

Think of them as the components for “a blueprint to architect for yes” inside modern, AI-driven enterprises.

“I divide the AI problem into four parts. How do you provision AI? How do you connect AI? How do you consume AI? And how do you protect AI in runtime? Your security strategy needs to map to these pillars.”

– VP Global Security, IT services, India

“The CISO mandate is to make AI visible, observable, governed, enforceable, auditable and defensible.”

– CISO, IT services, US



Visibility

A common mantra among CISOs is that they can't protect what they can't see. Visibility is a non-negotiable for effective security. Yet, according to Netskope data, an overwhelming 94% of enterprises report visibility gaps in their AI deployments, and only 9% can prevent a risky AI-driven action before it executes. Security leaders need to find ways to remedy this issue, tackling their shadow AI problem and improving their insight into vendor tools and supply chains.

Action 1

Start by removing the incentive to hide. Employees reaching for unsanctioned tools are usually trying to do their jobs, and treating them as adversaries only encourages concealment rather than compliance. The practical answer is to enable a authorized path that is genuinely usable, ideally consolidated onto a small number of platforms, with enforcement placed at the data layer. That way personal use remains possible, while corporate data cannot leave.

Action 2

Make the resulting visibility useful to the business, not just security teams. One organization issues monthly AI usage reports to a nominated ambassador in every team, helping them to understand what their own people are doing. Visibility becomes a service, and compliance follows.

Action 3

Accept that human review cannot keep pace. The volume of issues now being surfaced is far beyond what any team can triage manually, so automated detection is needed to help identify corrective actions at scale. As one CISO advised: "maybe it's actually AI that can help tackle the AI security problem."

"If you haven't provided a safe ecosystem or a contained environment for your employees to play with AI, it is highly likely they're off using it on the side with their own personal accounts."

– CISO, software, Australia

"Instead of treating employees as adversaries, provide coaching and guidance so that people know what the right decisions are."

– Steve Riley, Field CTO, Netskope

Governance

The governance framework and its related policies are where the security function sets out its intent. But the expectation that colleagues will read the governance documents and follow them carefully is under strain today, when AI adoption is distributed across the business and occurring at speed.

Action 1

Get policy off paper and embedded in products. Guardrails and harnesses can now have policies built in, to ensure security by design and reduce reliance on human workers to follow governance standards. As one CISO puts it, much of what sits in a policy document "...can be turned into machine language and actually embedded in the ecosystem". Security teams can also create a route for exceptions, so that an engineer who needs to vary a rule for a specific workload can request it.

Action 2

Establish a decision-making body between the business and the CISO. A cross-functional council that spans legal, privacy, vendor management, IT and the C-suite ensures all requests can find an answer, with a documented owner. That can be bolstered with named executive ownership for material use cases and quarterly board reporting on inventory, approved usage, exceptions and blocked actions.

Action 3

Make the governed route the fastest route. By pre-approving a small set of enterprise platforms, so anything inside them needs acknowledgment rather than assessment, security teams can focus their review effort on what is genuinely new. CISOs can also standardize what a review consists of, whether that's a fixed AI assessment questionnaire, independent penetration testing of vendor claims, or an external benchmark such as ISO 42001. They can do this by retaining evidence of the risk acceptance decision.

"We've been terrible in security, for forever and a day, around defining a security framework with a bunch of policies and standards, and expecting everyone to read them and know how to follow them."

– CISO, software, Australia

"We've formed a steering group with people from legal, ops, security and IT to review AI initiatives. That means we get diverse opinions at a time when the regulatory and security landscape is changing. So we get to hear from everyone and make informed decisions."

– VP Global Security, IT services, India

Controls

Until recently, controls were designed for humans. Now, AI agents create new vulnerabilities at machine speed. To be effective, modern controls must intervene inline, rather than simply logging what happened after the fact.

Action 1

Enforce at the data layer. Placing the control where data moves is what makes enforcement both inline and permissive at the same time. Security teams no longer need to block applications, so an employee can keep using their preferred service, but crucially corporate data cannot escape. This removes the main reason people route around security in the first place.

Action 2

Set and enforce boundaries for agents. Agents' actions must be evaluated before they execute, with a trust boundary set so that their remits are tightly bounded. As with human workers, it needs to be clear which action agents can perform, and on what data. Secure access service edge (SASE) or security service edge (SSE) frameworks should also help bring MCP traffic within existing security policies.

Action 3

Assume your controls will drift, and test them like an attacker. Currently most AI-related incidents can be traced back to a control that was set once and never checked. So each service needs to be monitored continuously, with controls adapted as it changes. One CISO has added what she calls a shift right to the more familiar shift left approach, pushing her team to "...act like a threat actor when it comes to AI and really try and break stuff," with a standing rule that anything found is fixed before an initiative goes live.

"AI is here and it's going to stay. It's not an external threat that CISOs need to chase. It's a constant variable they need to manage over time."

– CTO, IT services, Israel



Data protection

Part of AI's challenge to security teams is that it overturns the traditional approach to handling data. The most sensitive corporate information is now loaded into AI models deliberately, and comes back reshaped (and potentially polluted with hallucinations or poisoning). Security leaders must enforce modern data protection more rigorously.

Action 1

Approve the data, not the tool. Categorize your data first, then authorize each AI service against named data categories. That shifts the decision from whether a tool is permitted to what it's allowed to receive. Done well, this speeds up approvals, because a new use case for an approved service becomes a data question rather than a fresh review.

Action 2

Map data lineage before you trust output. CISOs should map data flows feeding each service, so they know where that data comes from and that it can be relied upon. This gives you greater confidence in AI outputs and the decisions made on the strength of a model's answer.

Action 3

Put retention and reuse in the contract; then verify it. Specify what a provider may retain, for how long, and whether your inputs may be used for training. Require independent testing rather than accepting self verification, and hold subcontracted fourth parties to identical obligations. These will need checking on a firm schedule, given how fast provider terms are changing. One security leader said he's seeing privacy policies at major providers change on a monthly basis.

"AI models require current, real data and struggle with synthetic data. They can create synthetic data sets for you, but it still means putting in your original crown jewels, so far as data is concerned."

– CISO, Australia

"Too many CISOs still frame their work around preventing intrusion when the real question for them today is preventing leakage through legit use. Because AI is legit. So it requires a conceptual shift, not just better tooling."

– CTO, IT services, Israel

Agent oversight

Security systems have been built on the assumption that every action has an identity behind it. From access rights to leaver processes, everything was attributed to a named human. But with the spread of AI agents, this approach breaks down. By the estimate of one CISO, in today's leading enterprises there can be 10 non-human identities for every human. This breaks down as AI agents gain greater access to collaboration tools, email, code repositories and identity systems.

Action 1

Give every agent its own identity, tied to the human who invoked it. Agents' credentials need to be tightly scoped, and ideally short-lived. CISOs should bind the agent identity to the person invoking it, so authorization can be evaluated as a single question about a specific human, agent, action and data. That blended identity combination helps bring agents inside the existing access model.

Action 2

Put agents through the joiner and leaver lifecycle. Agents are now being provisioned constantly but they are rarely decommissioned. Security teams must apply least privilege to agents, regularly review their entitlements, and terminate their access on the owner's final day, exactly as would happen with a normal user account.

Action 3

Supervise agents with agents, and record what they did. Human review cannot keep pace with autonomous activity. One option is to create security agents to monitor your other agents, acting as guardians over the wider environment. Each agent action should be logged so that incidents can be tracked and explained.

"The biggest thing many organizations are missing right now is identity for non-human actors."

– Steve Riley, Field CTO, Netskope

"It's a game of AI versus AI, with humans playing a critical role in the loop."

– VP Global Security, IT services, India

Confidence

If AI use is visible, governance is enforced, controls intervene inline, data is protected by category and agents are accountable, then approving a project becomes an evidence-based decision that CISOs can feel confident in. That's not the same as a risk-free approval, of course. But it ensures today's CISO can make recommendations to their CEO and board that increase business velocity and ultimately enable growth.

Action 1

Set minimum limits now and plan for continuous improvement in the future. AI is still a new and fast-evolving technology, so it's simply not possible to cover every eventuality right now. But CISOs can be clear about their baseline requirements for sensible security, while mapping out a plan to improve capabilities over time. As one CISO advises, it's important to get "really comfortable with being uncomfortable," yet commit to constant learning as this technology develops.

Action 2

Get the business to own the risk and put it in writing. CISOs should not be taking on the burden of risk acceptance for every tool and access request. Instead, they should present the risks plainly to the business stakeholder involved and ensure he or she understands the consequences. Putting the decision in writing provides confidence in the current process and future accountability.

Action 3

Report in business language. One of the best ways for the CISO to instill confidence in security decision-making is to give the board a small set of figures it can understand, framed within a compelling story and aligned to business goals. This could cover what has been approved, what has been blocked, what risk has been formally accepted and reduced, and what value that has enabled, perhaps on a quarterly rhythm. One experienced CISO urged his peers to "expand the skill set into softer skills like communication and storytelling."

"I can say yes to this project because I know that I've got visibility and control and governance in place."

– Steve Riley,
Field CTO, Netskope

3. Conclusion: Build before the gap widens

No technology in history has reached so many people so quickly as AI. It already pervades business conversations and boardroom meetings. It can be jarring to remember that, in its modern form, AI is barely four years old and we're still in the early stages of its development.

Yet even in that short time AI has created a plethora of new challenges for security leaders, from employees using personal tools, to agents roaming across datasets and systems.

To meet this moment, many CISOs are rethinking their approach and embracing an updated role as an architect of secure approval. They are building new protocols to speed up adoption, manage risk and defend their decisions.

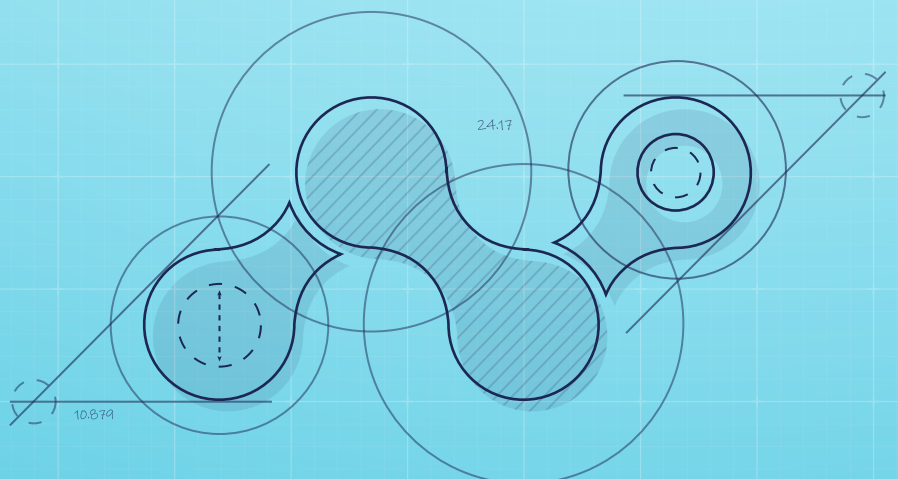
Now CISOs must accelerate that process, supported by expert partners, so they can become true Architects of Yes for their organizations.

“Start adopting new management routines and new ways of working, not just new tools... The question isn't what additional controls do I need, it's how do I operate differently now that part of my team isn't human anymore.”

– CTO, IT services, Israel

“You cannot govern what you cannot see, and most organizations are only governing a fraction of the AI actually in use. Given the speed AI is moving through industries, when the business asks to deploy AI, the CISO must be able to answer who decides, based on what evidence, and how fast?”

– James Robinson, CISO, Netskope



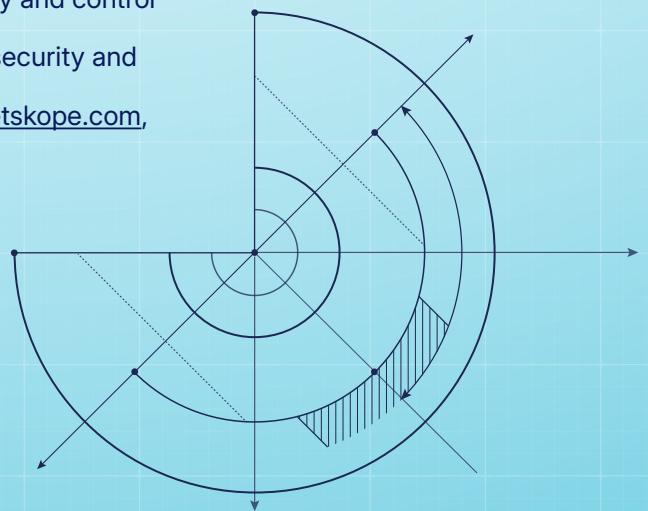
METHODOLOGY:

Phase one: Netskope conducted an AI security listening tour, meeting with 61 security leaders (CISO, CIO, CTO and VP level) across North America, EMEA, ANZ and APAC.

Phase two: deep dive follow-up interviews were conducted with 12 further leaders from five countries – the US, UK, Australia, India and Israel – during July and August 2026.

About Netskope

Netskope (NASDAQ: NTSK), a leader in modern security and networking for the cloud and AI era, addresses the needs of both security and networking teams by providing optimized access and real-time, context-based security for the AI ecosystem inclusive of agents, applications, tools, LLMs, people, devices, and data. Thousands of customers, including more than 30 of the Fortune 100, trust the Netskope One platform, its Zero Trust Engine, and its powerful NewEdge network to reduce risk and gain full visibility and control over cloud, AI, SaaS, web, and private applications, providing security and accelerating performance without trade-offs. Learn more at netskope.com, Netskope.ai, on [LinkedIn](#), and [Instagram](#).



©2026 Netskope, Inc. All rights reserved. Netskope, NewEdge, SkopeAI, and the stylized "N" logo are registered trademarks of Netskope, Inc. Netskope Active, Netskope Cloud XD, Netskope Discovery, Cloud Confidence Index, and SkopeSights are trademarks of Netskope, Inc. All other trademarks included are trademarks of their respective owners. 09/26 RS-1020-1