

Netskope and AI in the Fast Lane

July 23, 2026

By: [Frank Dickson](#), [Grace Trinidad](#)

IDC'S QUICK TAKE

[Netskope's](#) AI in the Fast Lane road show stop in New York City on July 14, 2026, showcased the vendor's effort to make AI-native operation governable, organizing Netskope One around a three-stage life cycle of discovering AI usage, securing the AI pipeline, and protecting runtime interactions, with AgentSkope agent risk scoring, MCP scoring, and the AI Command Center as that life cycle's operational layer. The event reinforces IDC's view that competitive differentiation is shifting from the underlying AI model to the harness that governs it, and IDC expects inline enforcement, AI acceleration, and transparent risk scoring to be the dimensions on which vendors in this category separate over the next 18–24 months.

EVENT HIGHLIGHTS

Netskope brought its global AI in the Fast Lane road show to New York City on July 14, 2026, drawing security and networking practitioners for a program on securing enterprise AI adoption. The stop was part of a worldwide series spanning North America, Latin America, EMEA, and Asia/Pacific and Japan, premised on enterprises having moved past AI experimentation into governance at scale.

A central theme was AI-native operation: treating AI as the default way of working rather than an occasional add-on, beyond merely having access to AI tools. Netskope organizes its AI security road map around four pillars: the AI-native platform, enabling and securing AI, accelerating AI performance, and AI agents for automation. Governance surfaced as a recurring thread, with responsibility for AI policy typically spanning InfoSec, GRC, HR, legal, and IT, and customer sentiment ranging from urgency over rising AI-related case volume to a preference for fast remediation over caution.

A second theme centered on Netskope One's AI Security pillar, positioned alongside Data Security, Secure Access, and Single-Vendor SASE, as the mechanism for governing AI use end to end. Netskope's 2026 *AI Risk and Readiness Report* found that 94% of respondents report gaps in AI activity visibility, 88% cannot distinguish personal from corporate AI accounts, and only 6% claim full visibility into their AI pipeline. On enforcement, 31% rely on written policy and employee compliance as their primary control, 11% report no controls at all, and just 23% enforce AI security in line at the point of action.

Netskope frames AI security as a three-stage life cycle. Discover AI Usage covers visibility into enterprise AI and MCP tools, AI coding assistants, embedded AI features in SaaS apps, and personal AI use. Secure the AI Pipeline covers risk assessment for MCP integrations, data classification ahead of model training, and AI red teaming for prompt injection and jailbreak vulnerabilities before deployment. Protect Runtime Interactions covers inline data loss prevention, jailbreak and prompt-injection blocking, and zero trust access extended to agentic AI and MCP traffic through an AI Gateway and Agentic Broker.

AgentSkoPe and the AI Command Center are the operational surface for the life cycle. AgentSkoPe attaches a risk score to enterprise AI agents and extends comparable scoring to MCP integrations, while the AI Gateway adds east-west visibility into server-side AI traffic. Enforcement is positioned as adaptive rather than binary: a policy engine fed by app, content, context, data, device, and user signals selects among coaching, encrypting, escalating, quarantining, and unsharing, including coach-and-pivot flows that intercept personal AI use and reroute the employee to a sanctioned corporate instance. CEO Sanjay Beri framed the strategy as "enabling AI is securing AI," describing AI as the operating system of work and citing model flexibility, data sovereignty, and performance, delivered through the NewEdge private cloud footprint with a combination of network peering and low-latency routing, branded as AI Fast Path, as Netskope's differentiators.

IDC'S POINT OF VIEW

In the view of IDC, enterprise security adoption continues to trail AI innovation by 18–24 months, a pattern IDC has observed across earlier technology cycles and expects to persist with AI guardrails and inline data loss prevention, both of which remain in single-digit adoption industry-wide despite years of vendor investment in the underlying controls. The gap traces less to a shortage of tooling than to organizational friction: Budget approval cycles and committee reviews move at a pace AI adoption has already outrun. IDC sees an opening for security teams to reposition around this reality, building governed on-ramps for AI use rather than defaulting to blocking, a shift that turns security from a bottleneck into a driver of adoption velocity. The market is consolidating around exactly this control point: within roughly the past year, SentinelOne acquired Prompt Security, Check Point acquired Lakera, Cato Networks acquired Aim Security, and Akamai completed its acquisition of LayerX, a wave that treats inline, prompt-side control of employee AI use as a platform requirement rather than a feature. Read against Netskope's own finding that just 23% of organizations enforce AI security inline, the acquisition wave is a leading indicator of where competitive pressure in this category will concentrate through 2027.

IDC also sees competitive differentiation moving away from the underlying AI model and toward the harness. Organizations now leverage that harness to enforce policy, manage cost, and limit exposure of proprietary data to frontier model providers. That shift favors vendors that can operate across many models, including smaller, purpose-built models that are often cheaper and more accurate than general-purpose frontier large language models (LLMs) for narrow, well-defined tasks. Netskope's own fleet of 191 proprietary models, only 20–30 of which are large language models, illustrates the pattern: Most of the work runs on smaller deep learning and shallow models, with LLMs reserved for cases that genuinely require them.

The scores such a harness produces warrant scrutiny of their own. IDC's cross-domain review of scoring systems in cybersecurity and AI security finds a market that has retained the vocabulary of confidence and risk while abandoning the statistical discipline behind it, blending confidence, severity, and likelihood into single numbers until meaning leaks out. For example, Netskope publicly details its calibration methodologies and reweighting rules within its Knowledge Portal. Netskope's risk scores are dynamic and continuously updated. Buyers evaluating the risk and MCP scores should ask all vendors how a score is calculated, what ground truth it is calibrated against, how fresh it is (how often scores are updated), whether it can be reweighted to their environment, and who sets the thresholds. A detection layer that recognizes only learned, known risks makes those questions more pressing, not less, and

vendors like Netskope and others that publish calibration evidence for their AI risk scores will set the benchmark against which competing scores are measured in RFPs.

An extremely interesting concern organically arose among the attendees. The frontier model providers can use what they learn from serving a customer to eventually compete with that customer. The risk is no longer hypothetical: Anthropic's April 2026 launch of Claude Design, a prototyping and presentation tool built on Claude, moved the company from foundation model supplier into the application layer occupied by two of its own design-tooling partners, Figma and Canva. Figma had integrated Claude models throughout its own product only months earlier; days after the launch, the Anthropic executive who sat on Figma's board resigned, and Figma's stock fell on the news. Canva reached a different outcome, converting its relationship into a multiyear partnership that threads Claude Design into its own workflow rather than replacing it. That split outcome is the risk IDC hears voiced most often in enterprise conversations: a frontier model provider that has spent months learning a partner's workflow can choose to compete with it, complement it, or land somewhere in between, and the customer has little say in which path it takes.

This shift also raises data quality to a first-order concern. Training and fine-tuning on stale or unsanitized data undermines model accuracy in ways that often surface only after deployment, making data security posture management and lineage tracking a precondition for safe AI use rather than an afterthought. It also exposes an unresolved governance gap: Enforcing intent and preventing drift depends on ontologies that most enterprises are not yet equipped to build or maintain, an early-stage challenge the market has yet to solve.

IDC flags AI acceleration as an underappreciated stage of this shift. Beyond simply removing blockers, acceleration extends to the technical performance of AI interactions themselves, including caching techniques that avoid repeatedly sending the same context and queries to external models, cutting both latency and token cost. Few customers are asking about this dimension of the market today, which suggests vendors have room to lead rather than react on performance messaging.

On positioning, IDC's guidance to vendors in this category is to resist framing AI security as a wholly new product category. Enterprises are better served, and more likely to buy, when AI governance is presented as an extension of existing data protection, threat protection, and access control capabilities applied to a new set of interactions, not a parallel stack requiring separate budget and tooling. Vendors that can demonstrate unified visibility across endpoint, browser, network, and inline AI traffic, activated within a platform customers already operate, hold a more durable advantage than those competing purely on model layer features.

Subscriptions Covered:

[AI Security and Trust](#)

Please contact the IDC Hotline at 800.343.4952, ext.7988 (or +1.508.988.7988) or sales@idc.com for information on applying the price of this document toward the purchase of an IDC or Industry Insights service or for information on additional copies or Web rights. Visit us on the Web at www.idc.com. To view a list of IDC offices worldwide, visit www.idc.com/offices. Copyright 2026 IDC. Reproduction is forbidden unless authorized. All rights reserved.