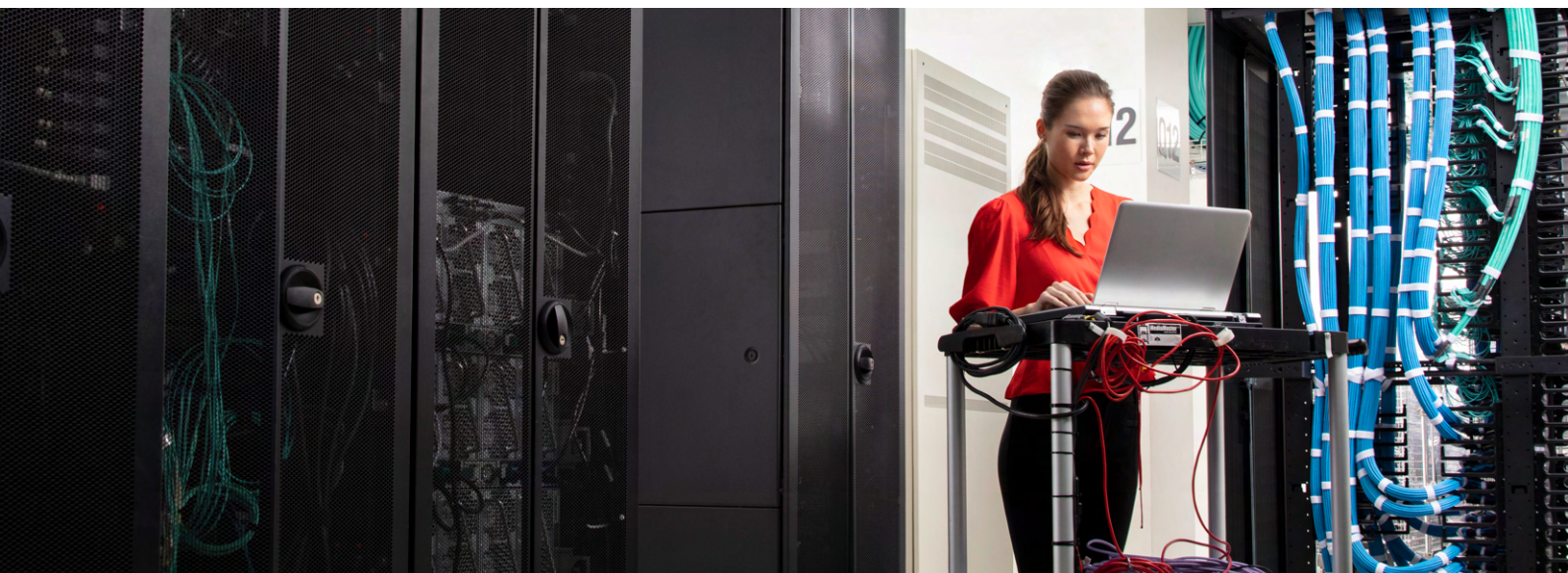




Building secure, AI-native networks: Four imperatives for the agentic era



AI is evolving fast, and your network has to keep up. How? Prioritize four imperatives to help ensure your setup can deliver performance, flexibility, and security—at scale.

At a glance

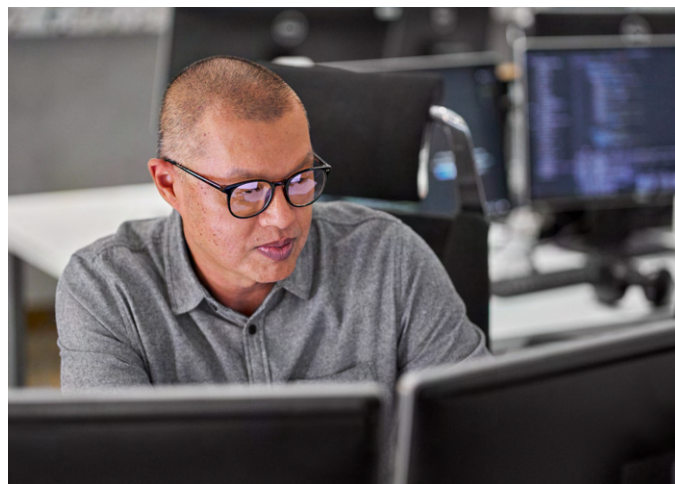
- Today's IT leaders are confident that their network is set up to support the increasing demands of AI, however, less than half admit to fully understanding the foundational network needs across the full AI lifecycle. This raises concerns about their readiness for the more complex demands of agentic AI.
- To rise to the challenge of agentic AI and beyond, the network must deliver secure AI-native operations and validated AI-optimized performance. In this new era, the network fills a more strategic role, actively transforming how AI is deployed, connected, and protected across the enterprise.

It's only been a year since HPE released its [last report](#) on the network's role as an AI enabler. But in that short time, AI has leapt forward—and we've entered a new frontier: **agentic AI**.

What does this mean for the network? The need for flexible, foundational support across AI's lifecycle—from data acquisition to model training and inferencing—remains. But we must now raise our expectations: the network must step into a more strategic role where it's actively shaping how AI is deployed, connected, and secured across the enterprise. **The question is: has enterprise understanding of this shift in role evolved at the same pace as AI?**

Our latest survey, detailed in [One year on: Architect an AI Advantage](#), reveals the answer. At first glance, confidence appears high: **87%** of IT leaders believe their organization's current network can handle increased AI traffic. But dig deeper, and gaps emerge. **Less than half** fully understand how the network can support critical AI stages like model development, training, tuning, or inferencing. And at most, only **54%** have a great understanding of data acquisition and preparation.

This disconnect between perception versus reality is a dangerous sign of overconfidence in foundational network support. And it comes at a time when the stakes have never been higher. Agentic AI adoption is already soaring: nearly half of respondents (**48%**) are actively deploying it, and another **42%** are planning to explore it in the next 12 months. These advanced AI models will demand far more from the network—and those demands must be met for continued AI success.



Your network's new role

The shift from broad support across the whole AI lifecycle to targeted operational transformation and readiness in the age of agentic AI indicates it's time for a new approach.

To unleash the necessary transformation, the new strategic approach to network design must prioritize:

- **Secure AI-native operations**, which is evolving from traditional “integral security” to agile, client-to-cloud automation built with AI for AI.
- **Validated AI-optimized performance**, where high performance connectivity is refined through workload-aware, validated optimization.
- **Sustainability at scale** that focuses on creating adaptive, energy-efficient networks for GPU scale at lower costs.

Focusing on these areas is the strategic key to modernizing your network for today's rapidly evolving AI needs. The goal is to have a network that is intelligently automated, secure by design, energy efficient and consistently high performing. It must be able to adapt to changing AI demands, whether in centralized data centers or distributed edge environments. And, finally, it must serve as a trusted foundation for deploying and scaling AI across the enterprise, supporting both operational efficiency and business outcomes.

“Without robust infrastructure and unified data governance, scaling AI beyond pilots is a pipe dream.”

- **ITOps executive in the public sector**, on the network's role in AI success



Network necessities for now and what's next

With the strategic approach outlined, let's explore the **tactical imperatives** for achieving these network goals. These four imperatives highlight key design and operational requirements.

Imperative 1: Support increasingly symmetric traffic

Networks must enable faster, more balanced data movement to support the growing complexity and real-time demands of symmetric traffic brought about by immersive AI applications. And this has implications for both data centers and user-facing access networks.

Traditionally, internet traffic has been asymmetric—we download far more than we upload. But increasingly **immersive AI apps have shifted this pattern. They need symmetric traffic support**, where uploading is as vital as downloading. As a result, access networks must now be equipped to **handle much higher bandwidth for two-way data**, and this requires upgrading uplinks from 10 Gigabits per second (Gbps) to 100+ Gbps. In data centers, this capacity will need to scale from 400 Gbps to 800 Gbps and soon to 1.6 Tbps.

Imperative 2: Deliver low latency in distributed architectures

Does AI have immediate access to the data it needs? As we enter the next frontier, the answer must be yes, because greater AI immersion demands greater data accessibility. To achieve this, architecture is becoming increasingly distributed.

- **Distributed data:** Once bound to a single location, data now spans multiple data centers and edge sites that sit closer to users and devices. This expansion addresses growing power demands while improving resilience at the same time. But physical spread introduces latency risk. Ultra-fast connectivity—via low-latency links—is essential to make these resources act as one system, enabling AI to access and process data instantly for real-time tasks like inferencing and retrieval-augmented generation (RAG).
- **Distributed workloads:** AI inferencing is delay-sensitive and must happen in real time, ideally in under 100 milliseconds. As workloads move to the edge, businesses need more specialized infrastructure: compact, high performance, energy-efficient networking designed for distributed AI. This enables lightning-fast responses and reduces inefficient practices like prompt recycling, which wastes GPU resources.

The architectural shift toward “distributed” also impacts data center-to-data center interconnections (DCI). Speeds must rise dramatically: think upgrading from local roads to high-speed motorways. Faster links mean smoother AI performance across locations. In short, AI now demands more responsive data center connections to support real-time, distributed processing.

But distributed doesn't replace centralized. In fact, it depends on it. A strong core network is the foundation for scaling distributed AI. Without a robust centralized data center network (DCN), adding edge nodes or interconnecting multiple data centers will create bottlenecks, reliability issues, and poor GPU utilization. Here, our survey findings revealed a gap in data center readiness: **only 45% have a strong centralized DCN**, making it impossible to deliver the low-latency AI experiences that businesses increasingly expect.

There are solutions that are purpose-built for the shift toward distributed. HPE AI data center technology is such a solution, providing seamless support for distributed AI training across data centers. It combines high performance, scalable, energy-efficient infrastructure with AI-native intelligence, with features capable of keeping GPUs synchronized and analytics running in real-time. And there's a strong business case for such a bridging solution: it unravels operational complexity, cutting costs by up to 85%¹ and reducing trouble tickets by 90%.²

Imperative 3: Build power-efficiency for GPU scalability

It's no secret that GPUs consume large amounts of energy. And this is set to expand dramatically in data centers running AI servers. Global data center power use is forecast to rise 165% by 2030. Analysts also believe that AI servers will scale from 8 GPUs to 576 GPUs per rack by 2027³—for context, this increase would draw the same amount of energy used to power 500 average U.S. homes. The takeaway is that **we can't take our foot off the power-efficiency pedal: we need GPUs to scale and maximize our AI ROI.**

¹ Financial benefits of Juniper Networks Wired, Wireless, and SD-WAN driven by Mist AI in managed network services, ACG Research 2022

² [ServiceNow customer case study](#)

³ Goldman Sachs warns AI bubble could burst datacenter boom, The Register, 2 September 2025.
Available at: [theregister.com/2025/09/02/goldman_sachs_ai_datacenters](https://www.theregister.com/2025/09/02/goldman_sachs_ai_datacenters) (Accessed: 19 November 2025)

GPU scaling is currently constrained by the individual data center's power envelope. AI workload distribution across multiple sites is the workaround. But this is where we need to pay attention: distributing workloads without maximizing per-site GPU density will backfire, putting you in a position where expensive infrastructure sits idle. Even small AI data centers with fewer than 500 GPUs can cost millions to deploy, making **GPU efficiency per kilowatt essential for protecting ROI**. Once again, high performance, energy-efficient networking is critical, because it's the only way to preserve power for core GPU operations. Fortunately, **53%** of our survey participants are using energy-efficient hardware and networking equipment to reduce the power consumption of their AI efforts. Our hope is that this figure increases as more embrace immersive AI within their enterprise.

“Our data center power budget is already maxed out—adding AI workloads is a challenge.”

– **ITOps executive in the energy / mining / utilities sector**, on capacity constraints of current data center infrastructure

Imperative 4: Secure AI agents and interactions at scale

AI-driven attacks from external threats are growing more sophisticated. And, in the agentic AI era, so, too, are internal threats. Unlike conventional AI systems, agentic models tackle tasks with more autonomy. This independence creates a new category of what analysts describe as “digital insider”⁴ threats, where agents may unintentionally (or maliciously) act outside of the guardrails of your enterprise's intended boundaries. Clearly, this poses an internal-threat risk that can rival or even surpass that of human insiders.

So, what role does the network have to play in helping AI-led defenses keep up with these escalating risks?

As we've seen, greater volumes of symmetric traffic generated by immersive AI demand smarter, faster infrastructure: a high-speed, low-latency network. And while traditionally the concepts of network **speed and latency** had only been linked to organizational performance, today they are also **important security features**. If the network is slow or congested, AI tools can't process traffic fast enough to detect threats or enable response in real-time to stop them.

This understanding of the link between security and network performance features is forming slowly. Among our respondents, AI-enhanced network monitoring tops the list of security tools, with a small but significant **40%** using this to protect the data feeding AI models.

As agentic AI deployments grow, networks must meet these performance and security needs **across a now vastly expanded web of traffic**—where complex agent-to-agent, agent-to-LLM, and agent-to-tool communications are creating more simultaneous connections and multidirectional data flows. This complex sprawl is driving adoption of technologies like **agentic mesh**, which delivers both automated traffic engineering and enhanced protection thanks to in-built features like real-time security monitoring, governed autonomy and resilient design.

Bringing strategic networks to life

Delivering on today and tomorrow's network needs means more than incremental upgrades. It calls for a fundamental shift in how networks operate: enterprises need smarter, more secure support. HPE's vision is a “**self-driving network**”. To this end, we're pioneering agentic AI-driven operations across our portfolio. The features and capabilities meet the needs of the agentic AI era:

- Using a common microservices-based architecture, Mist AI and HPE Aruba Networking Central now share architectural DNA, enabling a “develop once, deploy everywhere” model with technology sharing and cross-pollination of innovations that are **flexible, scalable, and interoperable**.
- Our platforms deliver autonomous AIOps through intelligent agents that proactively manage and optimize network performance across wired, wireless, WAN, and data center domains. These agents collaborate to troubleshoot in real time, automate remediation, and continuously learn to **reduce IT complexity and enhance reliability**.
- Extending HPE Aruba Networking Central agentic mesh with Mist AI introduces a multi-agent orchestrator, enabling AI models to reason together, coordinate actions, and **dynamically adapt to evolving network conditions**.

⁴ McKinsey & Company (2024) “Deploying agentic AI with safety and security: A playbook for technology leaders.”

Available at: mckinsey.com/capabilities/risk-and-resilience/our-insights/deploying-agentic-ai-with-safety-and-security-a-playbook-for-technology-leaders (Accessed: 20 November 2025).

Together, HPE's platforms create a **resilient, intelligent network** foundation that accelerates AI innovation, reduces complexity, and keeps enterprises secure at scale. We're helping enterprises elevate the network into its new strategic role, so they can move from AI experimentation to AI advantage—gaining speed, resilience, and differentiation in a competitive landscape that's increasingly shaped by real-time, autonomous intelligence.

About HPE

HPE is a leader in essential enterprise technology, bringing together the power of AI, cloud, and networking to help organizations achieve more. As pioneers of possibility, our innovation and expertise advance the way people live and work. We empower our customers across industries to optimize operational performance, transform data into foresight, and maximize their impact. Unlock your boldest ambitions, with HPE. Discover more at [HPE.com](https://hpe.com).

Learn more at

[HPE.com/networking](https://hpe.com/networking)



Visit [HPE.com](https://hpe.com)

[Chat now](#)

© Copyright 2025 Hewlett Packard Enterprise Development LP. The information contained herein is subject to change without notice. The only warranties Hewlett Packard Enterprise products and services are set forth in the express warranty statements accompanying such products and services. Nothing herein should be construed as constituting an additional warranty. Hewlett Packard Enterprise shall not be liable for technical or editorial errors or omissions contained herein.

a00156004ENW

HEWLETT PACKARD ENTERPRISE

hpe.com

