

LUMEN®

Wavelengths for AI:
**How Next-Generation
Optical Networks
Power Data Center and
Cloud Connectivity**



Custom content for Lumen from Studio by Informa TechTarget



The promise of AI is transformative—unlocking new efficiencies, insights, and competitive advantages—but most enterprise infrastructures are not prepared for the tidal wave of data and connectivity demands AI brings. AI-driven network traffic is on the verge of an unprecedented surge. By 2030, monthly global AI-generated traffic is projected to surge from 63 exabytes to over 1,200 exabytes, accounting for nearly two-thirds of all network activity worldwide.¹ This isn't just a matter of more data—it's a fundamental shift in how, where, and how fast data must move.

Almost nine in ten CIOs admit their infrastructure can't handle what's coming.² This is not a minor gap; it's a looming barrier that threatens to stall innovation at the very moment when speed and scale are most critical. "Enterprises are catching on very quickly, and without highly scalable network infrastructure, even the best AI architectures will underperform," says Todd Geneser, Director of Product Management, Lumen Technologies. Network infrastructure must keep pace before latency and cost become barriers to deployment. The bottom line is clear: For enterprises determined to lead in the age of AI, the network is no longer just a utility—it is the strategic backbone of innovation.



Almost nine in ten CIOs admit their infrastructure can't handle what's coming.²

1200 exabytes

By 2030, monthly global AI-generated traffic is projected to surge from 63 exabytes to over 1,200 exabytes, accounting for nearly two-thirds of all network activity worldwide.¹



The AI infrastructure shift — from Cloud 1.0 to Cloud 2.0

Cloud 1.0 had a simple premise: centralize computing in massive metropolitan hubs and connect everything via the general-purpose internet. While that worked for web applications, it isn't as well-suited to AI factories churning out industrial-scale AI.

Cloud 2.0 is a response to that. It's a shift in network architecture that merges cloud and enterprise core networks into a distributed, programmable, ultra-high-capacity network fabric. Network proximity now directly determines model performance. High-speed connectivity to those networks is a critical enabler for distributed AI architectures. It's a demanding idea because it requires rebuilding the roads while traffic continues to flow.



To make it happen, we'll need five things:



01

Extreme
bandwidth



02

Data center
interconnect as a
foundational element



03

Expansion into
AI corridors where
power is available



04

Distributed,
programmable
on-ramps



05

API-first network
architectures

Why AI breaks traditional networks

Traditional networks were designed for bursty, tolerant traffic such as email, web pages, and occasional video calls. AI workloads change the rules. Training an LLM model means continuously moving terabytes of data between storage and GPU clusters for hours or days at a time. It's a firehose that never stops, and smooth delivery without hiccups is non-negotiable.

Distributed training compounds the problem. Synchronizing calculations across dozens of AI nodes demands precision timing. One slow link degrades the entire job.

Connectivity is just as important as computing power in this environment, and CIOs who treat it otherwise are setting up their AI initiatives to stall.

The impact of bandwidth on data

10 PB	93 Days	9.3 Days	56 Hours
5 PB	46 Days	4.6 Days	27.8 Hours
1 PB	9 Days	22.2 Days	5.6 Hours
100 Tb	22.2 Hours	2.2 Hours	33.4 Minutes
Data vs BW	10 Gbps	100 Gbps	400 Gbps

The ability to move the massive amounts of data needed by AI further exposes the challenge that enterprises face on the road towards the digital evolution of their business.



The rise of the optical fabric — How wavelengths enable AI at scale

So if traditional network architectures can't deliver what AI demands, what can? The answer lies in dedicated, high-capacity optical channels carved out of fiber infrastructure, each one a private express lane on a multi-terabit highway. Wavelengths can move massive amounts of data at light speed between compute zones with none of the variability that plagues shared networks. Whether those clusters sit five miles apart or five hundred, the connection behaves identically with deterministic latency. It's fast, direct, and predictable.

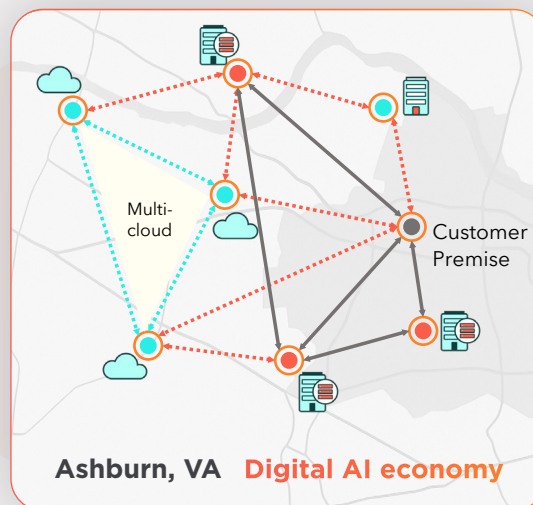
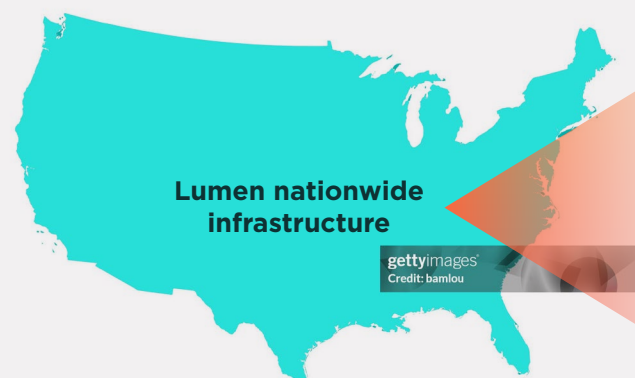
You can connect GPU clusters and IT instances across geographically dispersed data centers as if they were part of a single network. This is what people mean when they talk about an optical fabric, tying distributed infrastructure together into something coherent. It makes fast, smooth handoffs of AI traffic possible between data centers.

The Lumen Vision

Scaling connectivity to support the AI economy

2025

2026-2028+



Legend

- Data Center
- Hyperscaler
- Fabric ports
- Enterprise Premise
- Customer Premise
- Point-to-point connectivity (<100G)
- Lumen Connectivity Fabric (400G+ enabled)
- Point-to-point connectivity (100G+)

Metro + DC + Waves expansion:

Continued Metro and Data Center expansion along with new Waves Routes & capabilities

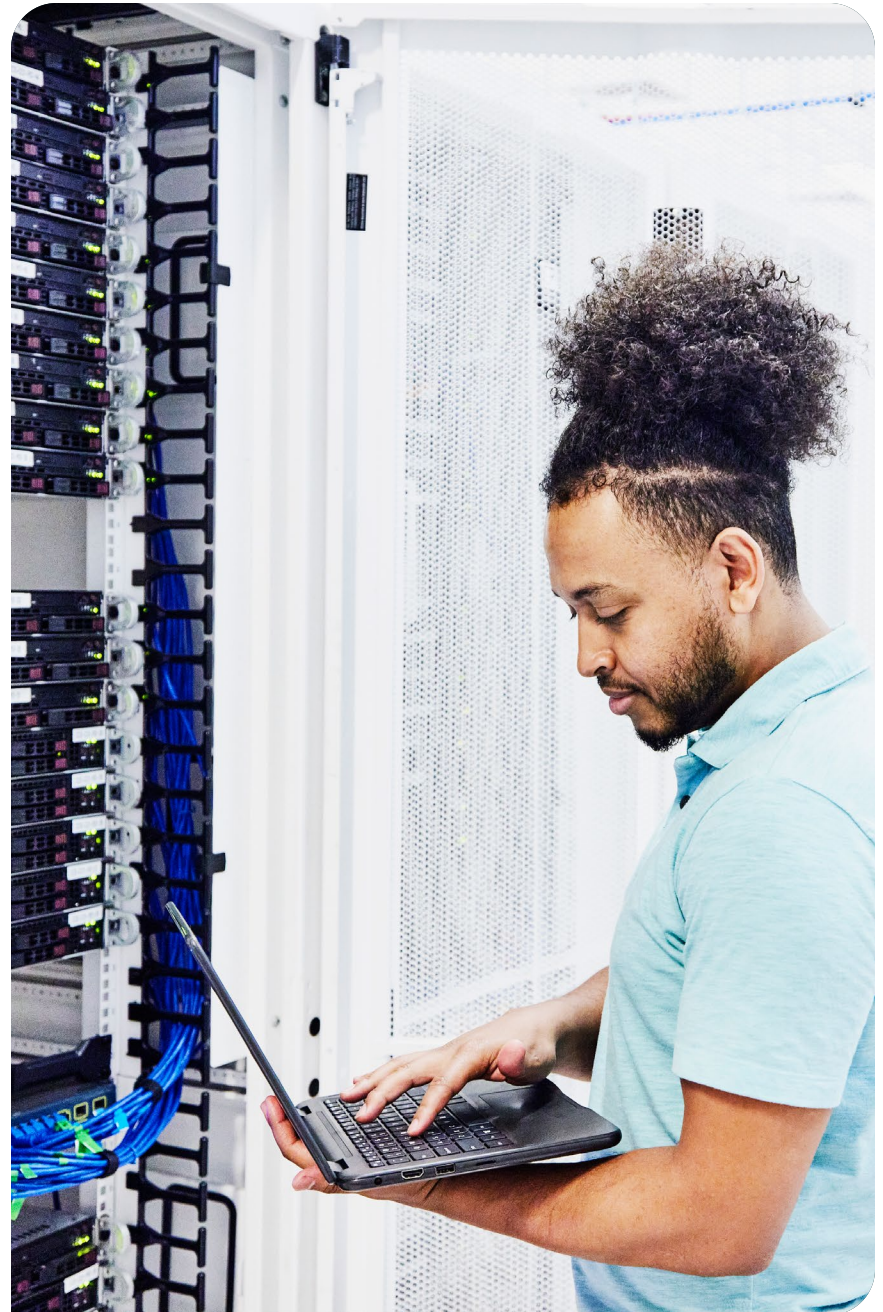
Direct, scalable access:

Bringing DCs and Hyperscalers “on-net” through high-capacity connectivity fabric, enabling multi-cloud workload management.

“Without dedicated wavelength connectivity, distributed AI architectures would be forced to leverage public networks. Dedicated network routes with terabits of available capacity act as the digital superhighways for the movement of massive amounts of data for training AI models and interconnecting with end user inferencing, enabling the continuous flow of data between the two.

Todd Geneser

The math works because of how optical networks scale. Stacking enough wavelength channels together using dense wavelength division multiplexing (DWDM) provides multi-terabit capacity, up to 50TB, on a single fiber strand. That's a strategic asset for organizations moving serious AI workloads today.



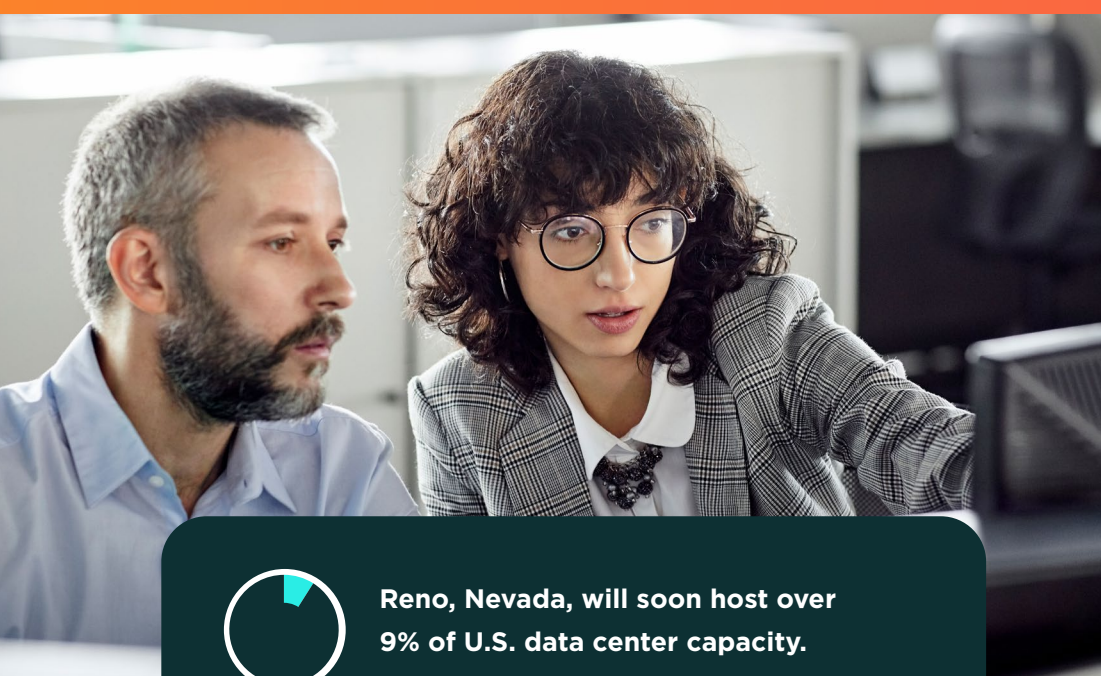
How data center geography is changing

For decades, data center strategy meant deploying IT instances at major interconnection hubs in densely populated metro areas such as Northern Virginia, Dallas, Chicago, and Silicon Valley. Ample power and network connectivity were the main factors, given the focus on Cloud 1.0 network architectures. Hyperscalers build data centers in these regions to enable the cloud on-ramp ecosystem customers need to more easily incorporate 'the cloud' into their overall IT infrastructure. Now AI is upending that model.

Training a large LLM model can consume tens of megawatts at a data center campus, and traditional data center hubs can't supply this scale of power quickly or cheaply. As a result, AI's growth is pulling the digital infrastructure into new geographies – wherever electricity is abundant, even if that's far from a traditional tech center.

“AI’s growth will be reliant upon the expansion of mega data centers into Tier 2 and 3 markets where the necessary power can be secured, Ample, reliable electricity, agnostic of form factor, heavily sways where the next big AI data center will be built.”

Todd Geneser

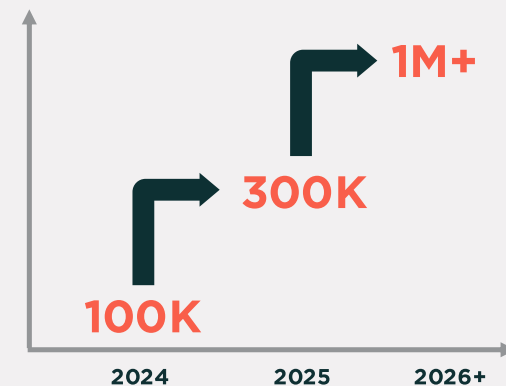


Reno, Nevada, will soon host over 9% of U.S. data center capacity.

Because of this existential shift, the industry is literally redrawing its map. Places once off the data center radar are now rising stars for capacity. Reno, Nevada, will soon host over 9% of U.S. data center capacity.³ Old power plants in rural Pennsylvania and Maryland are being converted into compute farms, while West Texas and North Dakota, previously known for cryptocurrency mining operations, are pivoting to accommodate AI demand. Hyperscalers are also appearing in other locations, such as Richland Parish, Louisiana. Canton, Ohio. Fort Wayne, Indiana. Montgomery, Alabama. These unlikely locales are being chosen based on energy capacity, not downtown convenience.

AI clusters beginning to drive more power than grids can deliver to a single site

GPUs for AI model training



Land

1M+

Size of data center campuses (sq. ft.) supporting AI infrastructure w/ expansion shifting to Tier 2-3 markets². Lower land costs and power availability potential



Power

8GW

Projected power for AI training run location by 2030³, driving model training to be geographically distributed



Fiber

3X+

Increase of fiber route miles to support future AI data center build by 2029⁴

Neoclouds are pushing this trend even harder. GPU-centric providers like CoreWeave (north of \$14.5 billion in funding), Crusoe (\$3.9 billion), Lambda (\$3.2 billion), and Nebius (\$1 billion) are purpose-built for AI workloads, offering a competitive alternative to the traditional hyperscalers.^{4,5,6,7}

CoreWeave is building a \$6 billion campus in Pennsylvania and investing another \$3.4 billion in UK expansion, while Crusoe operates renewable-powered facilities in Wyoming, Texas, Iceland, and the Nordic region.⁸

The numbers tell the story.

130 GW

Data center capacity demand is projected to hit 130 gigawatts by 2030, up from 41 GW today.

840 million sq ft

That requires roughly 840 million square feet of space, which is 600 million more than currently exists.⁹

This geographic shakeup to dispersed locations creates a new challenge: connectivity. Distributed AI workloads still need to communicate with each other at high speeds, so the network backbone must scale alongside the data centers. Fiber route miles will need to triple by 2029 just to keep up with planned builds.¹⁰ Providers are racing to lay fiber to rural and edge sites, and to upgrade backbones between new regional clusters. **The latest AI cloud geography demands a new class of interconnects: more direct, high-capacity, low-latency routes linking energy-oriented data centers.** The new backbone of digital infrastructure is optical, wavelength-based, and stretched between regions that barely existed on network maps five years ago.



Real-world applications: Where AI meets the network

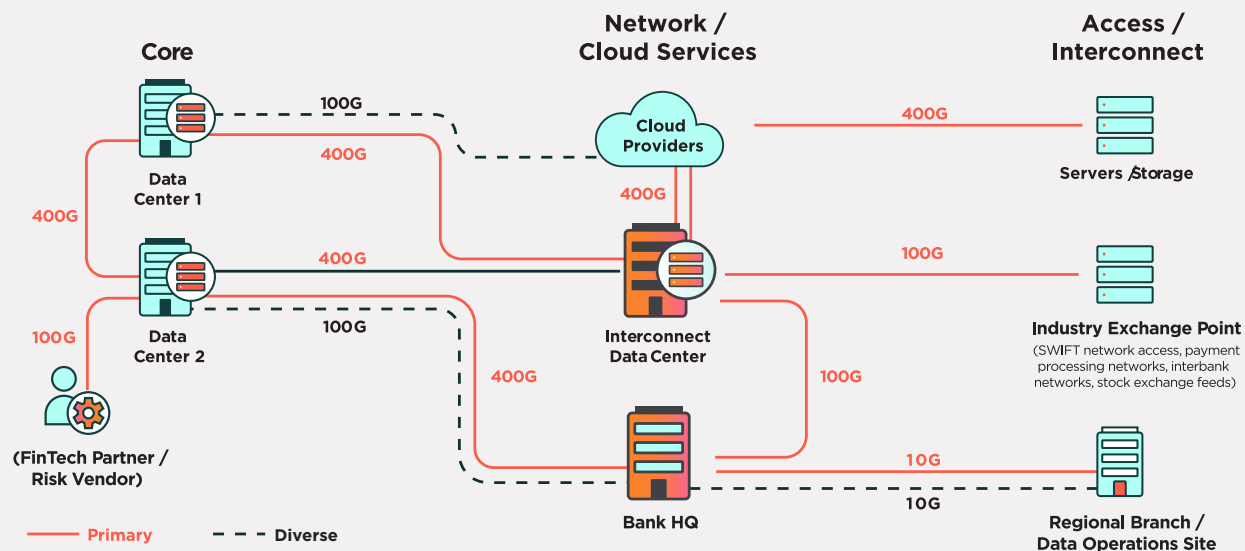
What does wavelength-enabled AI look like when it's actually running?

In financial services, downtime can be a regulatory event. Capital markets need maximum uptime and minimal delay for everything from real-time pricing feeds to stress-testing models that churn through terabytes of historical data.

One major bank deployed dedicated wavelength connections with AES-256 encryption and physically diverse routing between core data centers and cloud platforms. This enabled real-time market data flows and large dataset transfers for risk modeling.

Financial Services

From large real-time data flows to AI-powered risk models, financial institutions rely on low latency, encrypted connections across cloud, data center and core environments. Lumen® Wavelength Solutions provide dedicated optical transport with deterministic performance, built-in resiliency and the ability to scale to help customers support their compliance-driven workloads.



Healthcare presents different challenges but similar stakes. Hospitals and research networks move large imaging files, patient telemetry, and diagnostic data between facilities, data centers, and cloud instances that may be scattered across a region or country.

One national healthcare organization built a 100G mesh network that delivered zero downtime during a significant fiber cut. That's the kind of resilience you want when patient outcomes hang in the balance.

Media and entertainment seem like low-stakes environments. Still, a streaming provider delivering live coverage of a premiere event might think otherwise. Jitter and buffering at the worst time can tank subscriptions.

Major content platforms now use wavelength transport for high-definition delivery, to help ensure consistent performance even during peak demand. Behind the scenes, studios rely on the same infrastructure for real-time rendering, remote collaboration on visual effects, and live broadcast workflows.

In the pharmaceutical and life sciences industries, datasets are truly enormous. Genome sequencing, clinical trial coordination, and molecular modeling research often span continents and involve sensitive data subject to strict regulatory oversight. Getting a compound to market faster can mean billions in revenue; getting the connectivity wrong can mean compliance violations and delayed trials.

In all of these scenarios, AI performance directly correlates with network quality. Wavelengths provide the deterministic capacity and latency that enable data to move as fast as ideas depend on it—turning what would otherwise be a collection of distributed systems into a unified, responsive fabric.

What's next — Toward the intelligent connectivity era

The optical networks underpinning today's AI deployments are evolving towards becoming more scalable and dynamic throughout the customer journey, from deployment to assurance. Expect wavelengths channels to continue to grow. 800G wavelengths are expected in late 2026, and work to develop 1.6TB capabilities is underway. Lumen and Ciena demonstrated what's possible in 2025, completing a 1.2 terabit wavelength trial across 3,050 kilometers of ultra-low-loss fiber. It's the longest non-regenerated signal ever recorded at that speed.

The bigger shift is architectural. As Cloud 2.0 matures, the line between compute and connectivity will blur. Wavelengths won't just carry data between facilities. They'll function as extensions of the data center fabric itself, dynamically allocating capacity the same way virtual machines spin up and down today.

The market for this capacity is growing accordingly.

36%

The neocloud providers' GPU-as-a-service model will balloon from \$3.23 billion in 2023 to nearly \$50 billion by 2032, a 36% compound annual growth rate.¹¹ Each of those GPU clusters requires connectivity that can keep up, which means sustained demand for wavelengths infrastructure at levels we haven't seen before.



The upside for enterprises is flexibility. Wavelengths allow for quickly spinning up ultra high capacity connectivity as training sets and models become increasingly larger, allowing the network to operate with agility and in alignment with compute resources. They no longer have to wait months for circuits while competitors innovate.

This isn't just an infrastructure upgrade. It's a redefinition of how innovation moves: on intelligent, dynamic optical networks explicitly built for what AI demands.

The final perspective — Building the fabric of AI

Lumen has been assembling the physical infrastructure and service ecosystem that enables distributed AI. The foundation is fiber.

16.6 million miles

We have built the largest ultra-low-loss intercity fiber network in North America, with 16.6 million fiber miles already deployed and 47 million planned by the end of 2028.

60% more capacity

The network averages over 500 fiber strands per route, delivering roughly 60% more capacity than legacy infrastructure and enabling up to 125 400G wavelengths per fiber pair.



Raw capacity is only part of the equation. Speed of deployment matters at least as much, given how fast AI projects move from pilot to production. The Lumen® RapidRoutesSM initiative directly addresses this with 48 TB of pre-engineered capacity on 30 key strategic network corridors, enabling 100G and 400G connections with a 20-day delivery SLA to over 300 third party data centers, eliminating the custom engineering delays that typically stretch wavelength provisioning into months-long ordeals.¹²

The 400G buildout is aggressive. Lumen has expanded its enabled footprint to 100,000+ wavelength route miles, adding 1.39 petabytes of capacity across high demand intercity network routes. Connectivity extends to more than 125 cloud on-ramps with major providers including AWS, Azure, Google Cloud, Oracle, and IBM. We have access to 2200+ on-net third party data centers.

For latency-sensitive workloads, the network is designed to deliver sub-5-millisecond edge performance for 97% of U.S. business demand.

The goal isn't just more fiber in the ground. It's converting static capacity into programmable optical connectivity fabric. For organizations serious about AI, that's the difference between keeping pace and falling behind.

[Learn more](#) about how Lumen Wavelength Solutions can help your business transition to an AI future. →



Resources

1. Omdia, “Road to 2030: AI and the future of Network services – Traffic Outlook and Implications”, April 2024
2. IDC, “Enterprise Horizons 2024”, June 2024
3. Lumen, “Cloud 2.0: Because AI won’t run on yesterday’s internet”, October 22, 2025
4. Yahoo Finance/Reuters, “CoreWeave inks \$11.9 billion contract with OpenAI ahead of IPO”, March 10, 2025.
<https://ca.finance.yahoo.com/news/exclusive-coreweave-strikes-12-billion-182359737.html>
5. Crunchbase.com
6. Ibid.
7. Ibid.
8. ABI Research, “Profiling Six Leading Neocloud Companies”, October 27 2025.
<https://www.abiresearch.com/blog/leading-neocloud-companies>
9. SemiAnalysis. Multi-Datacenter Training: Open AI’s Ambitious Plan to Beat Google’s Infrastructure. September 4, 2024 (<https://newsletter.semianalysis.com/p/multi-datacenter-training-openais>)
10. Business Wire. FBA Research Reveals Critical Need for Fiber in AI and Data Center Growth. July 31, 2025.
(<https://secure.businesswire.com/news/home/20250731062095/en/Fiber-Broadband-Association-Research-Reveals-Critical-Need-for-Fiber-in-AI-and-Data-Center-Growth>)
11. Network World, “Neoclouds roll in, challenge hyperscalers for AI workloads”. June 25, 2025.
<https://www.networkworld.com/article/4011187/neoclouds-roll-in-challenge-hyperscale-rs-for-ai-workloads.html>
12. 20 business days starts from date Order & Order Addendum fully executed & received by Lumen



Lumen connects the world. We are igniting business growth by connecting people, data, and applications – quickly, securely, and effortlessly. Everything we do at Lumen takes advantage of our network strength. From metro connectivity to long-haul data transport to our edge cloud, security, and managed service capabilities, we meet our customers' needs today and as they build for tomorrow.

**For news and insights visit news.lumen.com, LinkedIn: [/lumentechologies](https://www.linkedin.com/company/lumentech),
Twitter: [@lumentechco](https://twitter.com/lumentechco), Facebook: [/lumentechologies](https://www.facebook.com/lumentechologies),
Instagram: [@lumentechologies](https://www.instagram.com/lumentechologies), and YouTube: [/lumentechologies](https://www.youtube.com/lumentechologies).**

Contact us

This content is provided for informational purposes only and may require additional research and substantiation by the end user. In addition, the information is provided "as is" without any warranty or condition of any kind, either express or implied. Use of this information is at the end user's own risk. Lumen does not warrant that the information will meet the end user's requirements or that the implementation or usage of this information will result in the desired outcome of the end user. All third-party company and product or service names referenced in this article are for identification purposes only and do not imply endorsement or affiliation with Lumen. This document represents Lumen products and offerings as of the date of issue. Services not available everywhere. Lumen may change or cancel products and services or substitute similar products and services at its sole discretion without notice. ©2026 Lumen Technologies. All Rights Reserved.



Expert led. Impact driven.

Studio is Informa TechTarget's global content studio offering brands an ROI rich tool kit: Deep industry expertise, first-party audience insights, an editorial approach to brand storytelling, and targeted distribution capabilities. Our trusted in-house content marketers help brands power insights-fueled content programs that nurture prospects and customers from discovery through to purchase, connecting brand to demand.

[Learn more](#)