

Infrastructure built for specialized cloud providers

Deterministic connectivity that helps you scale GPU platforms faster, protect SLAs, and monetize capacity with confidence.

As AI platforms scale, the network becomes a critical part of how the business operates. Connectivity starts shaping performance consistency, customer trust, and how confidently you can grow

The Challenge

Often, the network becomes the limiting factor as GPU density increases. Latency variability, slow inter-metro turn-ups, and fragmented network paths introduce unpredictability into training runtimes, put SLAs at risk, and leave valuable GPU capacity sitting idle.

That's when the network stops being just infrastructure and starts acting as a control point - determining how quickly you can expand into new regions, how reliably performance scales, and how efficiently capacity turns into revenue.



"Prometheus Hyperscale's partnership with Lumen furthers our aligned missions of building the backbone of the AI economy. Lumen's future-ready fiber is the connective fabric we need to provide best-in-class connectivity and drive AI innovation."

– Trenton Thornock

Founder and CEO, Prometheus Hyperscale

The Opportunity

As GPU platforms grow, taking control of the network becomes essential. Predictable performance, faster expansion, and the ability to bring capacity online quickly all depend on it.

How Lumen can help

Deliver deterministic performance at scale

Lumen provides high-capacity, low-latency connectivity between GPU clusters, data centers, and metros. Traffic is performance-isolated and engineered for AI workloads, helping keep training and inference consistent as demand grows.

Protect platform SLAs and customer trust

Built-in resiliency, redundancy, and security reduce variability across network paths, to enable performance to hold up under load and SLAs remain enforceable as environments expand.

Maximize GPU utilization and unit economics

By reducing activation delays and network bottlenecks, Lumen helps minimize idle GPU capacity and align network investment with real usage, so deployed infrastructure can start to generate returns sooner.

Common use cases

- **Distributed AI training and inference:** Deterministic, low-latency connectivity between GPU clusters helps deliver consistent training runtimes and predictable inference performance, even as workloads scale across sites.
- **Rapid multi-metro expansion:** High-capacity inter-data-center and inter-metro connectivity support faster region launches and customer onboarding, to reduce time-to-activation as new capacity comes online.
- **High-volume AI data movement:** 400G transport enables efficient movement of massive datasets between locations, helping keep GPUs utilized and unit economics in check as demand increases.

Network built for AI scale

High-capacity connectivity where GPUs run

Bring bandwidth to where compute lives, without redesigning the network as demand grows.

- **Optical Capacity** (100G/400G) to support high-bandwidth connectivity between GPU clusters, data centers, and metros
- **Private Inter-DC Connectivity** for low-latency, deterministic traffic across distributed training and inference environments
- **Elastic Capacity Scaling** to better align bandwidth growth with GPU utilization and customer demand

Performance-isolated network services

Reduce variability by isolating AI traffic and engineering the network for consistent performance under load.

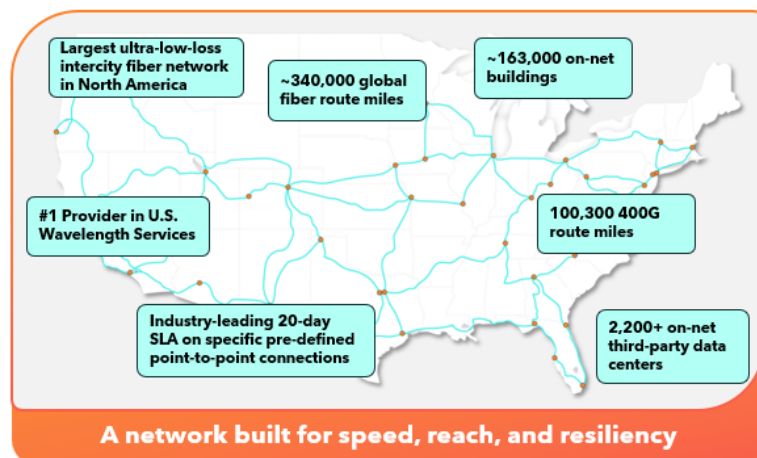
- **Dedicated, SLA-backed connectivity** designed to support GPU-intensive workloads
- **High-throughput transport** (up to 400 Gbps) optimized for large-scale AI data movement
- **Multi-path diversity** to support resilient, multi-data-center architectures
- **Built-in resiliency and security controls** to help reduce variability, protect data in motion, and support consistent SLA delivery

Architectural flexibility as platforms mature

Lumen offers flexible connectivity options designed to support GPU platform growth, helping you scale capacity, manage performance, and expand with greater confidence.

Solutions designed for your needs

- **Lumen® Wavelengths RapidRoutesSM**, for pre-engineered, validated capacity on select point-to-point routes, with as low as 20-day SLA on qualifying sites to help accelerate turn-up.
- **Lumen® Wavelength Services** provide high-capacity, deterministic east-west transport for GPU workloads, delivering the predictable performance required for multi-site training and inference.
- **Lumen® Managed Optical Fiber Networks (MOFN)** for the control and performance of fully dedicated Dark Fiber with Lumen-managed operations to extend private optical connectivity across data centers and cloud regions while reducing operational burden.
- **Network-as-a-Service (NaaS)** to support dynamic capacity needs and flexible consumption. Dynamic capacity with on-demand bandwidth changes support cloud-like consumption as traffic patterns shift.



Why Lumen?

Lumen delivers secure, high-performance connectivity built on one of the world's most expansive and deeply peered networks. With 30+ years of experience, we simplify complex networking challenges to help connect people, data, and applications across clouds, data centers and locations. Our continued investment in AI-ready infrastructure and metro expansion helps ensure your business is future-ready.

866-352-0291 | lumen.com | info@lumen.com

Services not available everywhere. Business customers only. Lumen may change, cancel or substitute products and services, or vary them by service area at its sole discretion without notice. ©2026 Lumen Technologies. All Rights Reserved.

LUMEN[®]
TECHNOLOGIES