

Large Enterprise

Backbone for AI:

Modernize the network to keep AI performance consistent



AI puts pressure on the network in two ways:

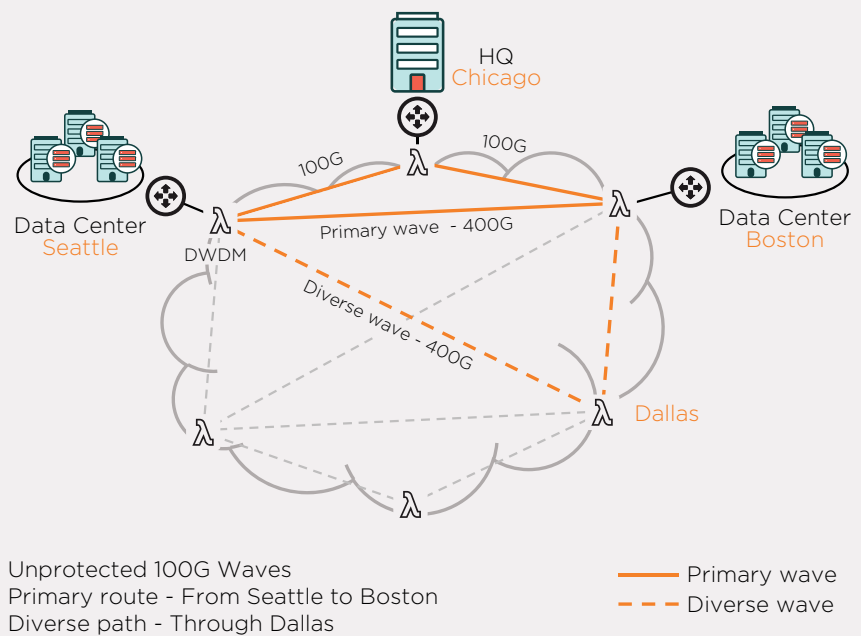
- **Training** moves large datasets and model artifacts across sites.
- **Inference** scales with adoption and becomes a persistent load. Performance issues show up as degraded experience, not always as outages.

AI traffic doesn't stay in one place—data, models, and users span sites, clouds, and edges. In large enterprises, that pressure shows up across locations, teams, and controls. Consistency matters as much as raw speed.

When the backbone isn't ready, AI initiatives slow down: longer iteration cycles, inconsistent app performance, and late-stage redesign.

Enterprise sites connect through edge/colo and cloud on-ramps, then out to users. The point isn't reach. It's predictable performance across the full path, including during maintenance and failures.

What the architecture shows



Decisions that matter early

1 Where inference must be served

- Which experiences can't tolerate variability?
- What latency range is acceptable for those experiences?

2 How traffic moves between sites

- How quickly can you add capacity when demand jumps—then jumps again?
- Which routes are long-haul inter-city vs metro, and what's your scaling plan for each?

3 Resiliency tier by business impact

- Where do you need speed to capacity (fast turn-ups, rapid upgrades)?
- Where do you need maximum control for sustained high volume?
- Where do hybrid managed options help reduce operational overhead without sacrificing performance characteristics?

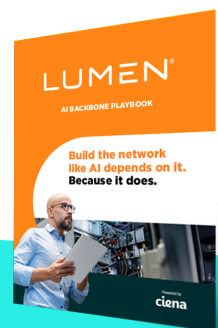
✓ Quick checklist (use in a planning meeting)

- ☐ Do we know which inference flows are experience-critical?
- ☐ Do we have true geographic diversity on those paths?
- ☐ If a route fails, do users see a slowdown—or a smooth experience?
- ☐ Can we add capacity quickly without redesign?
- ☐ Quick test: if a single construction cut can slow inference, you need true geographic diversity on that path.

What to do next

Map your inference paths and training flows, then choose connectivity and resiliency route-by-route (metro vs inter-city) based on performance impact—not habit.

Download the AI Backbone Playbook



About Lumen

Lumen is guided by our belief that humanity is at its best when technology advances the way we live and work. We deliver a fast, secure platform for applications and data to help businesses, government, and communities deliver amazing experiences.

Get in touch

lumen.com | 1-877-453-8353