

REPORT

Cloud 2.0: Because AI won't run on yesterday's internet

Uncover what's driving this shift and why enterprises and tech providers must adapt now to stay ahead.



Introduction

Cloud computing is at a turning point. For years, the first generation of cloud—what we call Cloud 1.0—shaped the digital landscape with centrally located, commoditized data centers and static network architectures. Today, the explosive growth of artificial intelligence (AI), surging data volumes and the proliferation of vast numbers of data centers in remote locations are transforming the digital landscape. These trends, combined with new demands for speed, security and flexibility, are driving the rapid emergence of Cloud 2.0—a new era that will fundamentally reshape how enterprises and technology providers operate.

Lumen partnered with leading consultants and analyst firms, leveraging proprietary data collection and market intelligence to map the future of cloud infrastructure. The findings are eye-opening, and draw on multi-source data, industry trends and real-world observations of data center and network evolution.

This white paper explores the drivers, dynamics and impacts of the monumental shift from Cloud 1.0 to Cloud 2.0. The research and analysis presented here will help you understand what's coming—and why you must act now to stay ahead.

What we thought of as cloud for 15 years is changing

Cloud 1.0 was an architecture centered on the commoditization of the traditional enterprise data center CPU, enabled by the hypervisor and associated platform and software services. It opportunistically used the telco interconnect network-to-network interfaces (NNIs) developed in the 1980s and 1990s and the public network infrastructure of the internet back when dial-up was relevant, the internet was emerging as a new communication superhighway—and software was shrink-wrapped and delivered on CD-ROM.

Yet Cloud 1.0 took decades to bloom, due to the entropy of the non-cloud world. Several factors influenced this slow pace, including the:

- Reluctance of IT departments to relinquish physical asset ownership
- Upskilling required for cloud literacy
- High cost of changing existing systems
- Reality that cloud infrastructure was never treated as a priority or considered an “easy button” for existing enterprise and provider networks

Cloud 1.0 emerged when Wide Area Networks (WANs) consisted of expensive, low-bandwidth static MPLS VPN meshes connecting all sites. This WAN architecture was eventually disrupted as Cloud 1.0 incorporated public dedicated internet access (DIA) and broadband internet links to add bandwidth for cloud applications. However, the Cloud 1.0 architecture wasn't built for, and didn't integrate well with, the migration to retail colocation and public cloud services. It simply used a bandage to help it cope with the strain of misalignment. That wrapper evolved into an infinite number of tunnel-over-the-internet solutions, never achieving interoperability or longevity. These challenges intensified with COVID-19 and the pressures of work-from-home migrations.

Today, Cloud 2.0 is emerging—driven by artificial intelligence (AI) and the resource-intensive interplay between GPU, CPU and storage. New tools, evolving communication patterns among Cloud 2.0 partners and distributed data storage, along with increased bandwidth and latency requirements, are also placing greater demands on the networks that interconnect these systems.

To address these challenges and demands, we'll need much more than an ace bandage and a reliance on the swamp of the public internet.

Figure 1

The elements of Cloud 2.0 that require a reset of network capabilities

Cloud 2.0 demands a network reset

AI is redefining the data center footprint, which is projected to experience 10X growth by 2030¹ to handle exabyte-scale demand

Legacy networks were designed for voice, internet and VPNs—**not bandwidth-heavy AI**

Cloud 1.0 connectivity with internet overlays and carrier-neutral facilities **has hit its limits**

Without purpose-built connectivity, billions spent on GPUs, SaaS and hyperscale **underdeliver**

Cloud 2.0 demands a reset because AI will overwhelm legacy network architecture

AI is not just another workload—it's a generational shift

Most analysts and consultants today are just now realizing the scope and significance of the change already in motion. In contrast, Lumen has made significant investments in data collection. We partnered with several of the few consultants and analyst firms that track data center and power development. The research-driven result is a vision that paints a clear picture of what cloud will become by 2030. Looking beyond the hype that AI will be huge, this vision maintains a focus on the impacts on Lumen and our customers.

Just a few short years into the AI boom, we see a shift in infrastructure to support the industrial AI that is taking shape and will rapidly unfold through 2030. Remarkably, many of these seismic shifts will occur by 2028.

This transformation includes a massive increase in use of dark fiber—mostly managed optical fiber networks—by some of the industry's biggest players, alongside alterations to the national fiber and power grids to connect their AI "factories"—purpose-built systems that continuously transform train new base models and retrain existing customized models for thousands of customers.

While we expect that these factories will primarily communicate with one another through distributed and sharded training and reinforcement at an industrial scale, they will also interface with the already overburdened Cloud 1.0 architecture in major metropolitan areas to distribute inference and exchange data and models with partners.

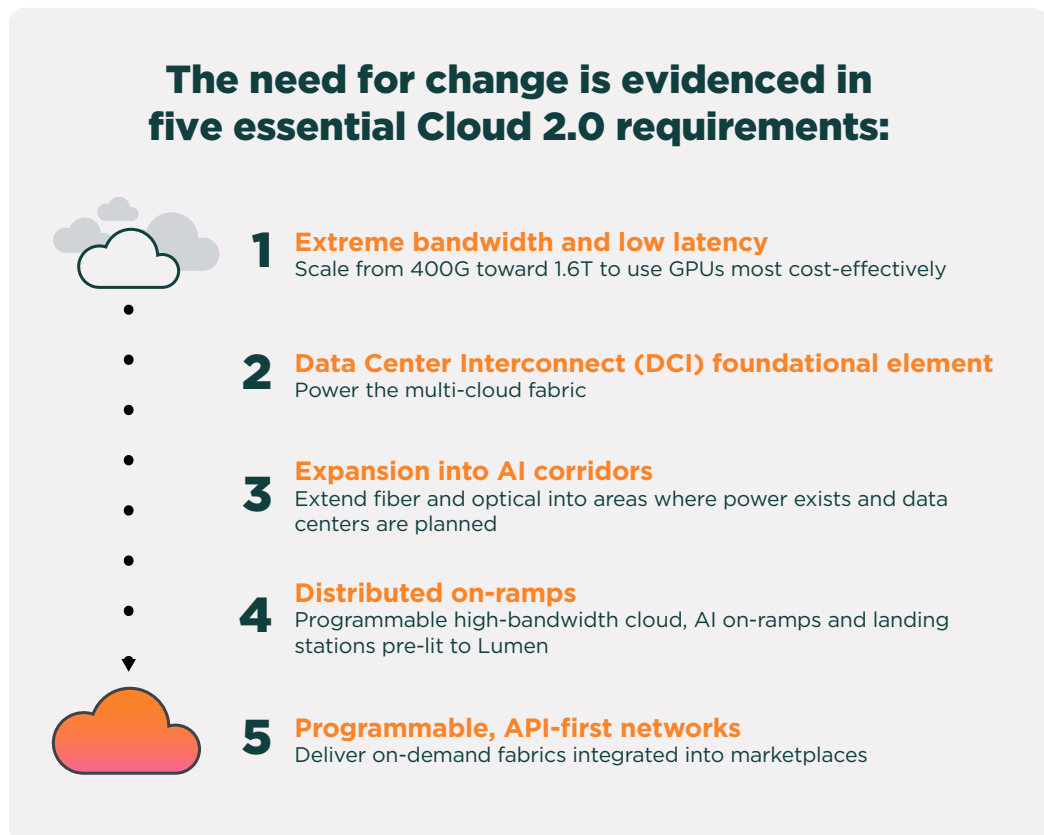
Resource constraints—including space, power and water supply—are pushing beyond the boundaries of these metro connections into more rural, less populated regions. Despite expectations of this tidal wave of growth, few analysts can articulate the necessary network requirements that follow.

Cloud 2.0: A transformation of cloud and enterprise core

Cloud 2.0 envisions a fundamental network architecture shift over the next 3-5 years, merging cloud and enterprise core networks to be faster, more secure, distributed and programmable.

Figure 2

Five new capabilities are needed to achieve a transformation of cloud and the enterprise core.



This once-in-a-generation shift will result in an independent and purpose-built network architecture strata. This is a significant departure from today's swamp-like internet, with its lack of bandwidth, redundancy and latency guarantees. Cloud 2.0 evolves beyond the outdated, static connection model that both providers and consumers of Cloud 1.0 were forced to use. Unlike its predecessor, Cloud 2.0 is designed to be programmable and flexible, meeting the demanding requirements that today's advanced services need.

Cloud 2.0 is ushering in flexible connectivity patterns over waves, Ethernet, IP and private IP with dedicated or burstable bandwidth, private IP integration and pay-as-you-go or flexible terms, with programmability to suit your specific needs.

A fundamental flaw in the belief that the public internet can be used for AI and Cloud 2.0 is that it's ubiquitous and infinite. The public internet may be ubiquitous, but it certainly isn't infinite. Just getting a bigger internet pipe will not be enough.

Driving massive expansion and structural change in AI infrastructure

Three key shifts in Cloud 2.0: **Densification, diversification and disaggregation**

According to research by Lumen and our partners, data centers will densify in select existing Tier 1 markets across all segments of builders—hyperscalers, wholesale and retail colocation builders, ecosystem primers, and emerging and neo-cloud specialists:

- Large concentrated markets like Northern Virginia will more than double power consumption—a proxy for density.
- Data centers in markets such as Phoenix, Chicago, Atlanta, Dallas, and Hillsboro, Oregon will double or triple in size.
- Space-constricted metro data centers with expensive real estate and power, such as New York, Silicon Valley and Los Angeles, will grow at slower rates, adding 10-20% additional capacity.

At the same time, data centers will diversify into adjacent suburban or rural locations:

- The region from Northern Virginia south through the Blue Ridge mountains into North Carolina will emerge as a data center corridor representing more than 9% of future U.S. data center resources (based on power) by 2030.
- New, more rural data center development is underway in the Midwest, Southwest, Deep South, Western Oregon, and eastern Oregon and Washington.
- Surprising hot spots like Reno, Nevada, will grow to provide more than 9% of U.S. data center capacity.
- Opportunistic conversions of power plants and old industrial sites with pre-existing power allocations are making rural Western Pennsylvania and Maryland home to big data center builds.
- Conversions of high-performance computing and cryptocurrency (HPC/Crypto) data centers are putting places like West Texas and North Dakota onto the data center map.

The most reliable segment of cloud projects—hyperscaler self-build locations—is appearing in Richland Parish, Louisiana; Canton, Ohio; Fort Wayne and Indianapolis, Indiana; Oklahoma City; Milwaukee; Susquehanna, Pennsylvania; Montgomery, Alabama; and Meridian, Mississippi, to name just a few locations that were not part of Cloud 1.0.

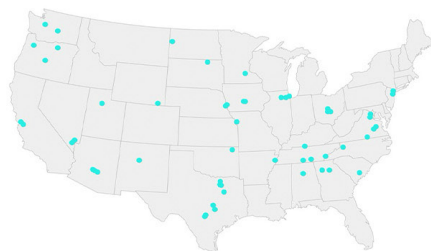
Hyperscalers are also developing or massively expanding locations in Tier 2 markets like Des Moines, Iowa; Reno, Nevada; Omaha, Nebraska; and Quincy, Washington, in addition to existing data centers in Tier 1 markets like Dallas and Columbus, Ohio. It's a breathtaking expansion of cloud infrastructure.

Figure 3

U.S. data center projects in 2024 and 2030 by location.

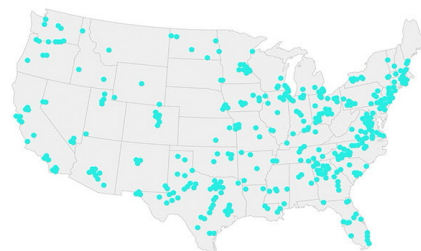
DC market all segments 2024

Data center projects¹ in 2024



DC market all segments 2030

Data center projects¹ in 2030



This explosion of growth is already driving a national (and global) network disaggregation as the demands of AI industrialization and AI-driven businesses require new stratified and purpose-built architecture that diverges from the general-purpose business and consumer, any-to-any internet architecture of the 1980s and 1990s.

Behind the numbers: Mind the gap

Whether you assume the U.S. data center market will quadruple (a more conservative Lumen estimate) or grow even larger (all announced projects), the outcome is the same: AI pushes data center growth sharply up and to the right—a significant structural infrastructure shift.

Figure 4

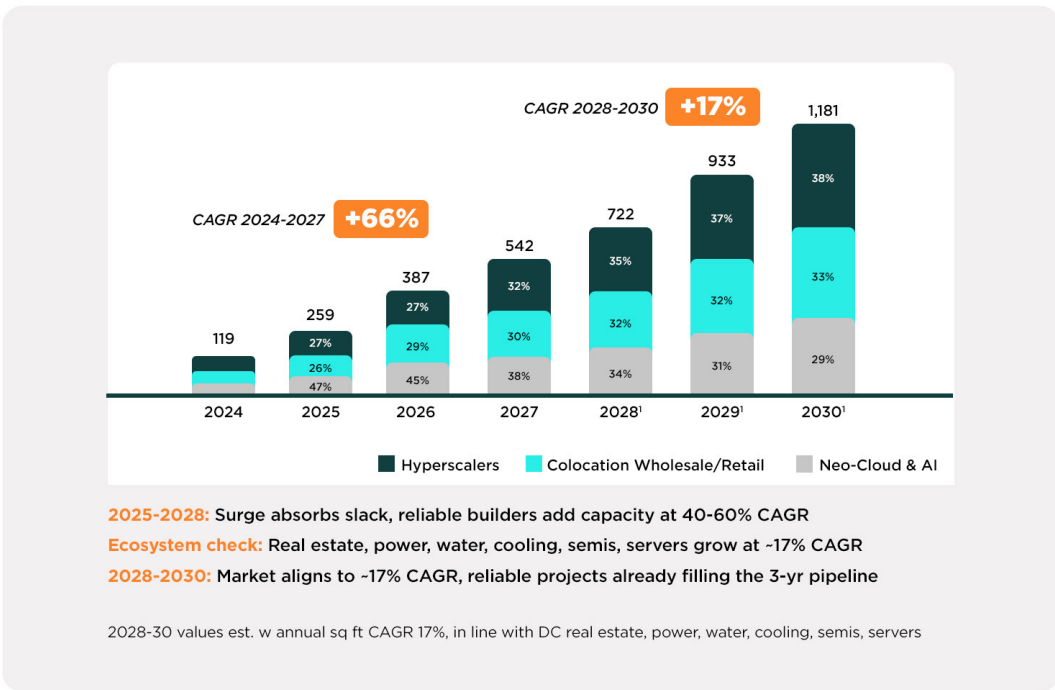
U.S. data center growth estimates from 2024 to 2030.



Lumen tracks all announced data center projects, and the data align with external market research—a jaw-dropping **~1 billion square feet** of incremental U.S. data center capacity added by 2030. To support this figure, Lumen examines the market through both a top-down view of announced data center build projects and a bottom-up view of build ecosystems, including real estate, power, water, semiconductors and servers.

Figure 5

U.S. cumulative data center growth estimates in millions of square feet from 2024-2030.



Announced projects

For the target segments—hyperscalers, wholesale and retail colocation builders, and emerging and neo-cloud specialists—Lumen analysis indicates incremental projects could peak in 2027 with a **68% compound annual growth rate (CAGR)** between 2024–2027. This is genuinely mind-bending growth, but one we should temper by considering identified ownership. Across all segments, this yields a very conservative Lumen estimate of **~400 million square feet** of additional space by 2028 (Figure 10 and Figure 11). It’s important to note that large data center projects typically take two to three years to complete, so visibility beyond 2028 is limited.

The build ecosystem view

Beyond the high confidence of a three-year data center build cycle, we can also test CAGR and investment cycle projections of the essential ecosystem components of a data center build, namely real estate, power, water, cooling systems, semiconductors and servers.² Across these components, the **weighted average growth rate is between 16-18%**.

However, a cycle mismatch is hidden in this rate, most notably in the 5 to 10-year development cycles for land, power and water.³ In other words, the three-year component-level growth projections and 5 to 10-year infrastructure development cycles conflict. So, to be conservative, we lean into data center build projections in areas with the fewest power and water constraints and anticipate data center expansion growth in constrained markets.

Figure 6
Development cycles and CAGR of market components.

Component	2025-26 projections	2028-30 outlook	Growth rate (CAGR)
Real estate	25% to 66% by 2030	75+GW total capacity	19-27%
Power provision	PJM: + 70GW, ERCOT: +67GW projected	80GW demand target	20%
Water resources	US DC needs 276B liters	1200B liters globally, 480B in the US (I)	11.5%
Cooling systems	Liquid cooling adoption 13% -> 85% penetration	\$63.6B market size	13.7%
Semiconductors	\$593B global spend	\$500B in AI Chips by 2028	15-20%
Servers	\$205B in AI Servers	\$384B in AI Servers	17%

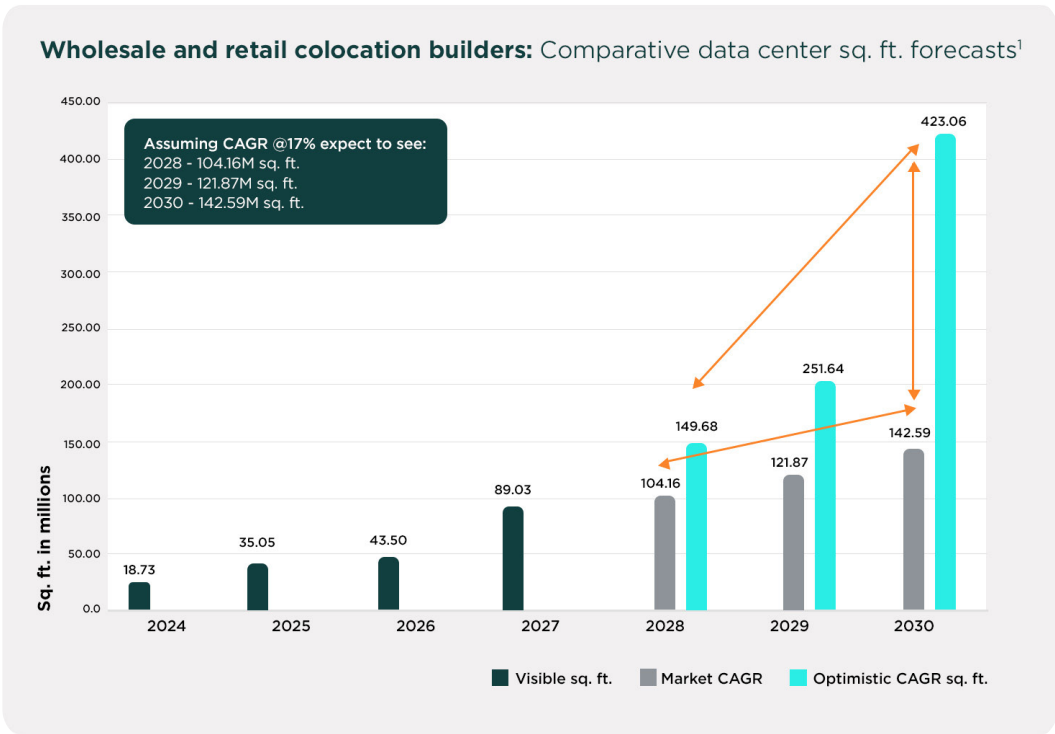
From the component CAGR, it is apparent that the current bulge in CAGR might create a false sense of systemic future capability. The growth boom through 2028 is consuming all the resource slack in the system.

A practical projection of data center expansion

If you use the average data center component CAGR for the latter part of the projection window (2028-30), you arrive at a more grounded projection for future data center growth. Figure 8 illustrates how this applies to the wholesale and retail colocation builder segment, highlighting the gap between market component capabilities and projections based on the current CAGR.¹

Figure 7

Adjusting the CAGR to market component averages leads to a potentially more reliable projection beyond 2028 for a segment of builders.



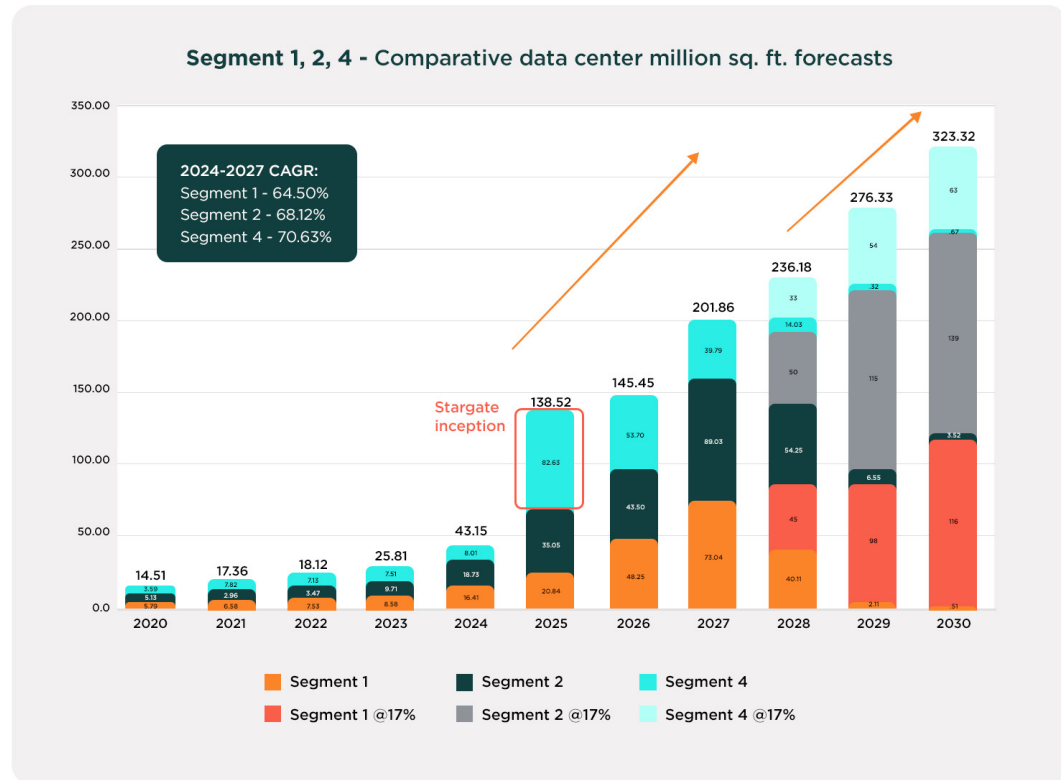
The market can potentially reach the 1 billion square foot mark across reliable data center segments in the 2025–2030 window (Figure 9). Here is why that is credible without over-rotating the near-term spike in growth. Approximately 120 million square feet of data center space was built between 2020 and 2025 across hyperscalers, wholesale and retail colocation builders, and emerging and neo-cloud specialists. During 2025-2027, the period where we have the strongest line of sight, the industry is on pace to add another 120 million square feet or more each year. In other words, we could see **400 million or more** square feet of additional data center space in just three years.

For the later phase, we examine how new data center developer projects may be absorbed under market resource-constrained CAGR. Our projection anticipates market resource support for another potential 800 million or more square feet by 2030 at adjusted growth rates. We continually track new announcements and conversions into identified tenancy.

As of August 2025, projects on record for 2028-2030 were sufficient to fill approximately 470 million square feet of that potential capacity—and we're just entering the higher confidence window for those builds. The velocity of these new projects is impressive—from May to August 2025 alone, newly- announced projects total approximately **100 million square feet** of potential additional capacity¹.

Figure 8

Adjusting the 2028 and beyond CAGR to average component CAGR across all reliable data center segments adds more than 1 billion square feet between 2025 and 2030.



As the Lumen analysis demonstrates, ongoing monitoring of announced projects and the relevant components is essential to maintain accurate projections and an understanding of the forces underpinning and shaping Cloud 2.0.

Goin' up the country: Where infrastructure growth is occurring

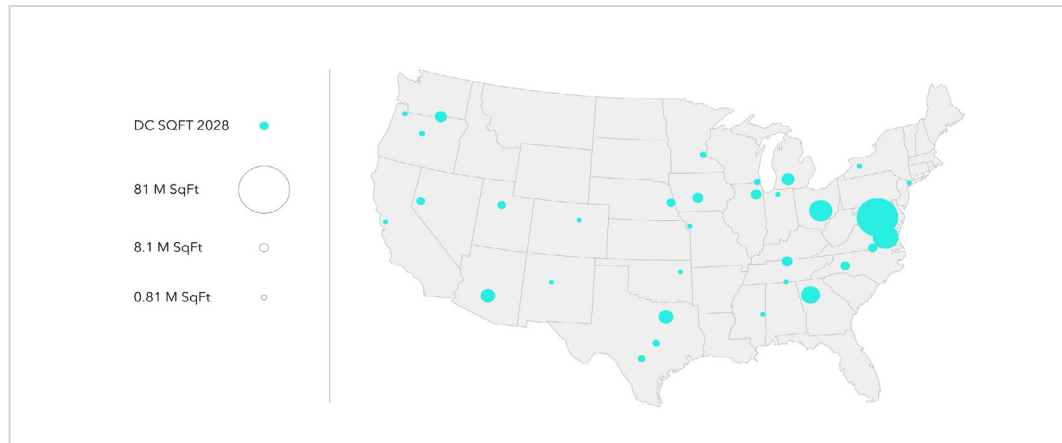
It's not only the hyperscale providers that are moving to the country. The whole cloud ecosystem is shifting.

New data center builds in remote areas involve all segments of cloud infrastructure builders. Hyperscalers and neocloud AI specialist projects combine to provide a high-confidence and conservative estimate of U.S. data center growth.

Lumen is tracking projects that will drive a fourfold expansion of the existing U.S. data center footprint—from approximately 98 million square feet to almost 400 million square feet by 2028 (with an additional 63 million square feet coming by 2030).

Figure 9

Relative scale of new data center space from projects built by high confidence sources through 2028—a fourfold increase in size with accompanying sprawl.

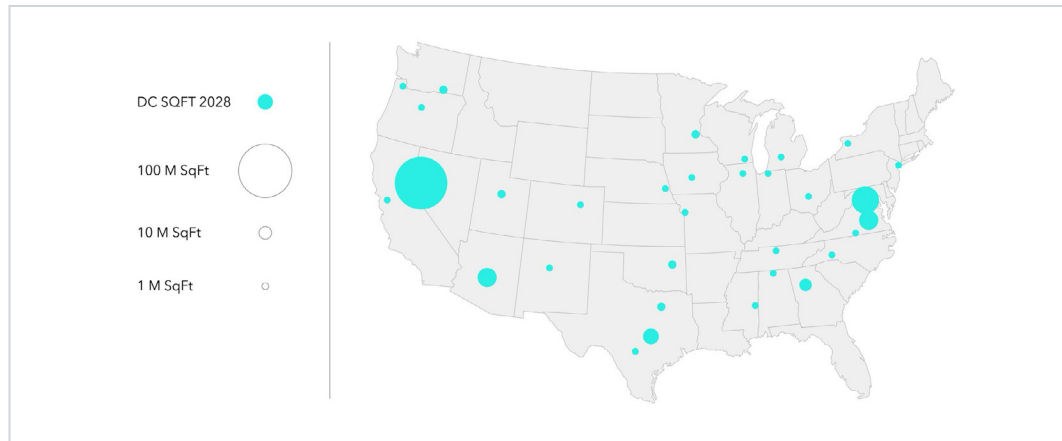


A slightly different view of where the growth is happening considers the density of proposed builds, bringing forward some of the notable builds mentioned earlier.

For example, those planned by emerging master planners and ecosystem primers are tracked with slightly lower confidence. These projects convert into higher-confidence builds as they progress; for example, by acquiring power and anchor tenant commitments. Such projects point to potentially enormous additional growth, adding approximately 600 million square feet by 2030 and potentially driving the overall available data center space in the U.S. to a 10X multiple of its existing footprint—an order of magnitude leap.¹

Figure 10

Relative scale of new data center projects by ecosystem primers magnifies the projects of other sectors and highlights new areas of the map—potentially moving Reno, Nevada into the same ballpark as Northern Virginia, the current world leader in data center space.



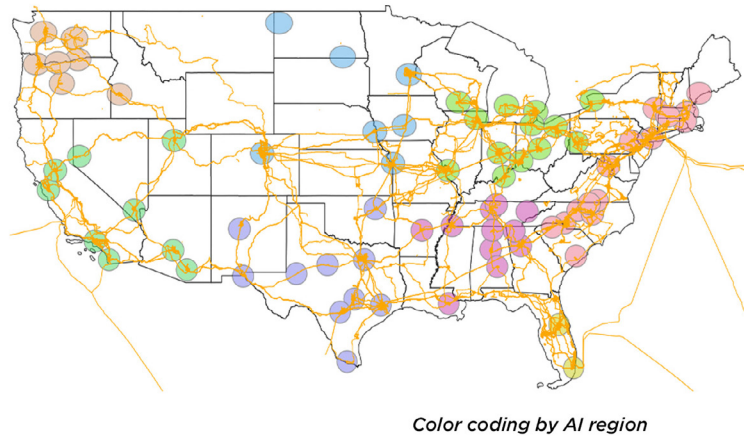
However, this explosion of rural data center operator (DCO) clusters only further exacerbates current architecture problems. Instead of backhauling to a major metro to find a carrier neutral facility (CNF) interconnect, new local interconnection may be much more efficient. Lumen has identified dozens of new data center clusters scattered across the U.S. that will require fiber, wave and IP services.

Figure 11

New data enter and hyperscale cloud clusters that offer opportunities for localized (often non-amplified) or optimized optical connections for DCI and IP.

Cloud 2.0 is redrawing the network map. New AI regions will emerge with new corridors.

Analysis of AI traffic and medoids



The implications are clear for data center providers and enterprises: *Your business is probably moving. Your workloads will be in different places.* With vacancies in existing major metro data centers at all-time lows, digitizing enterprises will need to migrate resources.

Are you planning a major project that will come online in the next couple of years and require five megawatts of power? Reserve your place now—but realize that it might not be located near your current operations. You can also consider partnering with a company that specializes in scouting, acquiring and developing large-scale data center parks. While you can hope for regional proximity for your workloads, it may be cheaper to run them in a greenfield data center in a different region.

Disaggregation: The end of flatness

The connection to and between data centers is now a crucial element of Cloud 2.0, moving away from the old notion of the enterprise WAN core. Relying on public network overlays is outdated. We are now building modern and programmable private networks without the constraints of the past.

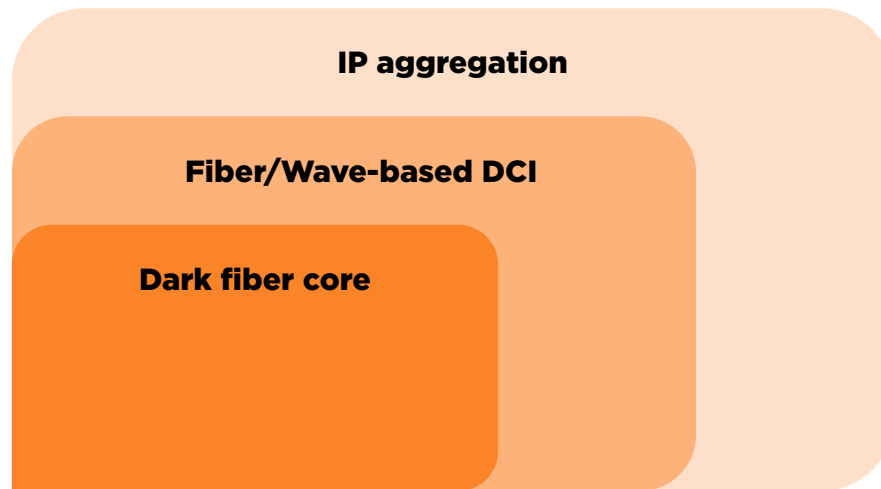
Data center interconnection is the new network core

Instead of connecting everything equally, a very distinct architectural stratification is emerging. The previous flat structure of enterprise networks is evolving into distinct, purpose-driven networks:

- Dark-fiber-enabled (often provider-specific and private) networks for industrial production and backend Cloud 2.0 applications
- Ethernet and waves-based DCI-heavy Cloud 2.0 front-end networks that connect applications, storage, AI inference engines and training
- Broadband and direct internet access (DIA) edge transaction networks that aggregate individual and branch access to cloud-edge SaaS for security and other cloud IT functions

Figure 12

The flat internet connectivity model that dominated the past 25 years of network architecture is breaking into three separate strata.



There is no longer a need for one “magic ring” to bind them all at the protocol layer, a principle that defined decades of a service provider-enabled multiservice backbone. Today’s networks are fabrics, and they are evolving, bound at their respective edges by L2 and L3 gateways that are likely API-and-cloud delivered.

Despite predictions extolling the growth of public network use via overlay technologies, networks beyond the access or edge strata will still be solidly private, high-bandwidth fabrics to control performance and security, and in the case of dark fiber, the asset itself. However, software-defined wide area networks (SD-WAN), Secure Access Service Edge (SASE) and other cloud-native networks will not go away. They are beautiful, simple and useful, and make life easier and faster.

The fundamental change will occur with commercial broadband circuits and “the internet swamp.” In the Cloud

2.0 AI economy, these will transform into a virtualized, programmable underlay that uses dedicated bandwidth from premises to the cloud and create fabrics in near-real-time as needed. Application workflows will be separated by SD-WAN and SASE tunnels, each transported over a virtual underlay that is controlled for bandwidth, latency and redundancy according to the specific requirements of each service.

In the Cloud 2.0 AI economy, SD-WAN and SASE will evolve into a virtualized, programmable underlay that uses dedicated bandwidth from premises to the cloud and create fabrics in near-real-time as needed.

It’s naive to imagine a telecom future that doesn’t digitize while all its customers do.

These private services will be digitally delivered, API-driven and integrated into hyperscaler marketplaces and enterprise workflows. This is an obvious yet difficult enterprise evolution and not a side project. Similarly, telcos will face a choice—be a true Cloud 2.0 provider or a consumer services company.

While the bandwidth increase requirement in core networks is obvious, the access network also needs a reboot, even as most companies are still grappling with zero-trust network, SD-WAN and SASE rollouts and migrations from DIY and bespoke managed service providers (MSPs) to edge cloud SaaS. Partially coupled to the unknowns around inference architecture, the only reasonable estimate of requirements is more API-controlled bandwidth services, not overlay tunnels. This applies to both the individual access links to and from those metros to the rest of the production fabric.

Low latency is yet another consideration. The current belief is that inference will be done at the edge, but with high-megawatt data centers within five milliseconds of much of the U.S. population, you have to question pushing inference to the power and space-constrained edge or shoehorning it into more resource-restricted on-premises spaces, particularly in fiber-rich geographies. Lumen research finds, for example, that manufacturing and drug research complexes will be built with separate on-site manufacturing, power and data center facilities. To achieve the extremely low latency needed for fully automated, “lights out” operations, the data center must be physically integrated into the facility’s design and directly connected to regional or national Cloud 2.0 sites.

It’s not just the concept of the edge in Cloud 2.0 that is different, but networking itself.

Agnostic data center interconnects are redefining connectivity

With the uncertainty of Cloud 2.0 network demand for industrial training, retraining and inference, and the emergence of data center interconnect fabrics as provider-based core architectures, it seems only natural that network and bandwidth consumption will become more flexible, even with a large fabric port size. An important element of this is that the data center interconnect is agnostic to the data center operator. Enterprises will choose data center partners without the concern of walled gardens or ecosystems. An island approach looks exceedingly small compared to the ocean that is Cloud 2.0.

New meet-me points

This evolution is already underway as part of the rush to fulfill Cloud 2.0 demands. Carrier neutral facilities (CNF) built in the 1980s as telephone exchanges were originally convenient for building early data center interconnects. But these CNFs were never suited for high-bandwidth DCI and are not needed as centers of traditional Internet architectures. They’re also physically highly concentrated; the entire nationwide on-ramp architecture for hyperscale cloud providers consists of approximately 120 on-ramps. Of the 51 distinct buildings that hold on-ramps, 17 contain three or more—comprising more than 40% of the entire cloud infrastructure.

Direct connectivity to public clouds typically launched about five years after the initial rollout of cloud services and was retrofitted into CNFs.

Cloud on-ramps have been squeezed into these spaces for convenience. But both the growing demand for direct connectivity and the expanded footprint of Cloud 2.0 are no longer a good fit for these sites.

For older facilities, as space and power become more costly and bandwidth demands rise, it makes more sense to move Cloud 2.0 on-ramps to dedicated regional transport hubs that connect directly to cloud provider regional network nodes. Scores of regional broadband and wireless providers don't factor into this part of the solution space, particularly when the minimum entry is 100 Gbps connection.

While CNFs previously valued having hundreds of providers for network interconnection, today only a handful provide the fiber necessary for 100-400 Gbps data center interconnect in any region and nationally.⁵

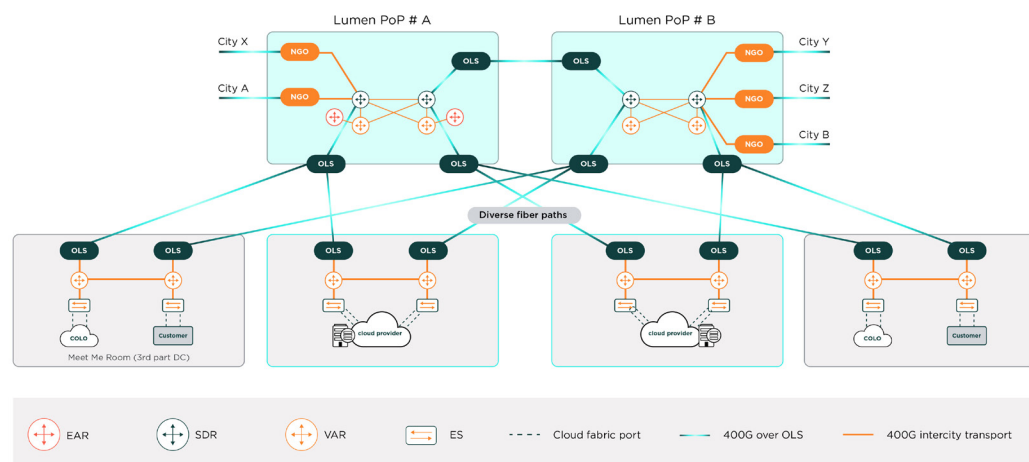
The old CNF meet-me-rooms (MMRs) may still be appropriate for building interconnects or NNIs, but not data center interconnects at scale—they are unsuited for Cloud 2.0 demands. Over-pulling fiber while consuming space and power for optical interconnection when only a few providers per metro are capable doesn't make any design sense. Interconnection must move into operator-neutral data centers and out of competitive commercial retail and wholesale data center space. An “extension cord” to Cloud 1.0 on-ramps will permit continued aggregated lower-speed connections.

Demand for multi-100G and 400G connections is making transport providers and hyperscalers rethink the CNF model. This has become painfully evident with the need to interconnect to new large rural DCO clusters at high speeds. Rather than building unnecessary backhaul to metro CNF interconnects, local interconnection may be much more efficient. The emerging architecture offloads CNFs via extension cords from metro points of presence (PoPs) to transport hubs designed with regional and national optical interconnection and 400G access directly from customers to hyperscaler networks, bypassing a third-party data center for interconnection.

Moving forward, Cloud 2.0 connectivity will focus on data center operator backbones, interconnect fabrics and new on-ramps wired directly into service provider and hyperscaler networks, potentially giving more life to existing CNFs via network-based extension cords.

Figure 13

A shift in cloud meet-me architecture to a transport optical basis more directly integrated with cloud provider networks, with extension cord(s) to existing MMR architecture.



Implications for the enterprise

How can enterprises transform to Cloud 2.0 without missing a step?

COVID-19 accelerated digitization using Cloud 1.0 resources, fundamentally changing how enterprise workplaces operate. Now, Cloud 2.0 will further distort the infrastructure of traditional business. In a way, the digitizing enterprise is approaching a point of identity that their born-in-cloud competitors started from that you are your workloads, and much less tethered to any physical place.

It's hard to say this shift is unexpected. Although we expected it to be more gradual? It may be time to move from talking about companies born-in-cloud to those with "AI in their DNA."

The resource demands of AI training, coupled with competition for resources, have significant implications for the data networks that move finished models and the vast amounts of private data used for training and retraining.

A recent RBC publication on AI notes that training has begun to involve exabyte-scale data transfers, with an eye-opening graphic on time required to transfer such data over various speed links.⁶ For example, to egress an exabyte of data over a 10 Gbps connection would take an astonishing 1,389 hours to transfer. When conducted over a 400 Mbps or 800 Gbps connection, the time decreases to 694 and 347 hours, respectively. The path to even higher speeds, such as 1600 Gbps and 3200 Gbps, is evident. This fundamental networking technology transition can't come fast enough, and the suppliers are behind.

You may think an exabyte transfer is extreme or an outlier. Still, even a petabyte transfer—a comprehensible size to a growing group of enterprises—can realize significant efficiencies through higher bandwidth connections. As illustrated in Figure 15 below, the cost of idle GPU time can be substantial at those sizes. *It's no wonder the industry is experiencing an inflection in 400G demand.*

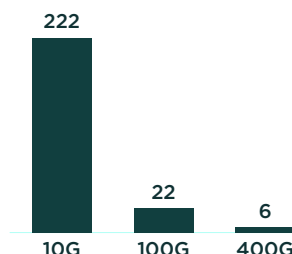
Figure 14

Cloud 2.0 requires enterprise networks built for exabyte-class moves; 400G+ is a starting point and the need for 800/1.6T is obvious. At petabyte and exabyte scale, the network, not GPUs, sets training time and cost.

Without bandwidth to DCs, GPUs sit idle and enterprises can't transform AI

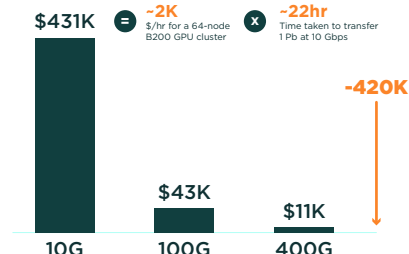
Enterprise network bandwidth has a significant impact on AI training cost...

Hours to egress 1 PB of training data



...and transferring 1 petabyte of data for training can reduce cost by ~40x with 400G vs. 10G¹

Enterprise Cost before training can start (\$K)



The daily delivery of AI inference could trigger new and distinct network effects. AI inference spans a spectrum of design options including model sizes, types and input/output interactions. These range from simple prompt-style chatbots to more complex implementations such as Retrieval-Augmented Generation (RAG), which integrates local data and uses search or query inputs as model inputs. It could be multi-modal (working with text, speech and video) or agentic, with sharded data and scoped agents or high-performance data queries in conjunction with agents. The network requirements for methods beyond simple prompts are demanding, and most industries realize that they must move beyond these to experience real value. What this means is that the trinity of API-driven bandwidth, latency and redundancy is critical to align with workloads and workload management.

Because of data privacy concerns, enterprises may choose to localize models, running them in private clouds to prevent data leakage upstream into provider models. Thus, the data hub and agent protocols of Model Context Protocol (MCP) and Agent2Agent (A2A)—which are in their infancy—become new policy control points beyond the “can we make it work” approach we are using today.

Interconnected enterprises will most certainly transform with Cloud 2.0. Both public and private internet networks must become faster—as well as more secure, distributed and programmable—to meet new demands. Most of the architectural foundations that defined Cloud 1.0 are obsolete and cannot support the requirements of the new cloud era.

Most of the architectural foundations that defined Cloud 1.0 are obsolete and cannot support the requirements of the new cloud era.

The best guess of analysts and providers alike regarding a Cloud 2.0 impact is an anticipated downstream effect that pulls more customers into their own managed optical fiber network (MOFN) or wave networks to connect their multiple data center footprints. Because Lumen is at the forefront of this growth, we are seeing a tremendous spike in demand for metro and long-haul dark fiber, as well as lit services—far beyond anything predicted by industry analysis services prior to 2025. Our own forecasts see a continual demand for years to come.

The Cloud 2.0 economy: accelerating innovation, connectivity and commercial change

A few years ago, it became evident that everything would be cloud-enabled via the enterprise digitization wave. And now it is apparent that eventually, everything will be AI augmented. Cloud 1.0 will connect to Cloud 2.0, and you may not even be able to see the seams. This will lead to a new Cloud 2.0 economy.

Discussions around how to handle storage for AI training and augmentation have made the concept of storage cloud meshing real. Security cloud is already here, along with about every IT service or solution delivered from a cloud. We’re going to need data and service delivery fabrics—clouds connecting to clouds.

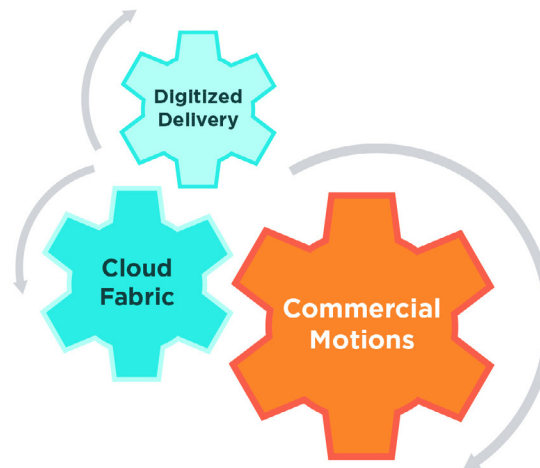
The multi-cloud nature of Cloud 2.0, the DCI mesh center and outer access IP edge will push product concepts like virtual layer-3 gateways—also called multi-cloud gateways or cloud routers—to bind them into new territories that merge public with private IP spaces, thereby enabling virtual extranets.

This is bound to have profound effects on channels and commercial motions. Naturally, consumption and delivery will take a step forward. All of it must be delivered from an API-first platform such as Network-as-a-Service (NaaS), but APIs need to integrate into DCO and hyperscale marketplaces. Also, Design, Price, Order (DPO) platforms must replace the Configure, Price, Quote (CPQ) telecom approach of the past.

The fabric port concept mentioned earlier—combined with the varying bandwidth demands of steady-state inference and periodic training or retraining—will require updates to the automation and tooling currently used in NaaS offerings such as Lumen Digital. These changes will also influence how the underlying networks are engineered and managed.

Figure 15

The Cloud 2.0 economy has implications for commercial motions, connectivity and service delivery.



What's next

Behind the curtain

Behind the curtain of the Cloud 2.0 economy lies more than network architecture segmentation.

Unlike digitization, which enabled enterprises and providers to make a smooth transition, AI-driven Cloud 2.0 is rapidly reshaping the landscape and introducing new techno-economic divisions.

On the enterprise side, the old categories for networking and infrastructure—large enterprises, mid-market businesses, and small and medium businesses—need to evolve into one of these new classifications: AI in the DNA, Using AI, Transforming and Laggard.

Both providers and consumers are moving much closer to a true cloud ecosystem where the network doesn't own the customer. The outcome is cloud.

On the provider side, even though the AI segment of Cloud 2.0 is seeing strong growth and a vibrant startup ecosystem, there is potential for further consolidation—resulting in a marketplace dominated by a few winners. Fundamentally, the greatest competitive advantage for providers in the AI marketplace is access to capital for infrastructure.

Predicting vs. preparing for the future

It doesn't take a crystal ball to predict that the second wave of an established technology will happen astonishingly fast—tech history is rife with examples. But this version 2.0 has the potential to break more than a few things.

The rapid pace of growth and opportunity in the cloud sector, combined with the significant capital required to build and support the necessary infrastructure, means that disciplined planning is essential. At Lumen, we're studying maps developed by our research teams that track, plot and evaluate the certainty of all the upcoming data center developments and all critical fiber routes that need to be created to supply the data. This view shapes Lumen infrastructure investments and directs conversations with existing hyperscalers, reliable data center build partners for wholesale and resale, and emergent AI infrastructure providers.

Lumen has already announced plans to grow our U.S. inter-city fiber footprint from 17 million fiber miles to 47 million fiber miles through 2028 as we also expand our metro fiber footprint. We're working closely with our cloud and data center operator partners to architect new paths to connect enterprises to Cloud 2.0.

Do you see what I see?

When I talk with customers, I see a strong overlap in vision between those who are more advanced in digitization and the early adopters of AI. These organizations already have data center-centric interconnects in their network topologies. They are also fully multi-cloud, often at bandwidths that push them into wavelengths or the higher reaches of shared Ethernet services in NaaS offerings .

We welcome anyone interested in the evolving techno-economics of our industry to join this conversation.

Figure 16

Enterprises must meet these key requirements for to successfully adopt Cloud 2.0.

Enterprises will be defined as leaders or laggards: Cloud 2.0 won't wait

Enterprises are going AI native: Cloud-first isn't enough. Inference, privacy and multimodal workflows demand expansion across multi-cloud. First movers define industries.

Networks need to change fast: Exabyte transfers and real-time inference will overwhelm today's overlays. Cloud 2.0 change is predicted to happen 10x faster than Cloud 1.0.

Programmable fabrics to be enabled: DCI-centric, on-demand connectivity flexes with AI workloads through self-serve, secure, zero-touch platforms.

Customer experience must be digitized: From static catalogs, manual provisioning to API marketplaces, self-serve fabrics, click-to-buy and zero-touch turn-up.

Do you hear what I hear?

In the coming months, I look forward to sharing Lumen research-informed roadmaps and talking to more of you about your existing topologies and the architectural and solution choices you're making for your Cloud 2.0 networks.

This paper began by highlighting the urgent need for reinvention to meet the demands of Cloud 2.0. Next, I'll explore the technical innovations and unresolved bandwidth control challenges that are essential for enabling Cloud 2.0—both on-premises and at its on-ramp points. These solutions support the broader architectural and operational shifts discussed earlier. In the context of evolving relationships between cloud providers and connectivity partners like Lumen, these innovations offer a path to once again “make smart pipes out of dumb pipes.”

About the author

Dave Ward



Dave Ward is Chief Technology Officer and Product Officer for Lumen. In this role, he is responsible for the development, integration and deployment of the Lumen global network, while driving the creation of disruptive technologies through innovative product development.

Dave has more than 25 years of experience reinventing and expanding the capabilities of networks and the value they bring to customers. He previously was the CEO of PacketFabric, where he oversaw its Network-as-a-Service (NaaS) offering and led significant product expansions. Prior to

that, he served as a Senior Vice President, CTO-Engineering and Chief Architect at Cisco Systems, and as a Fellow at both Cisco Systems and Juniper Networks.

Dave holds a bachelor's degree from Syracuse University and graduate degrees from the University of Minnesota.

Acknowledgments

This white paper was authored by Dave Ward, who led the research, writing and overall development.

Dave gratefully acknowledges the contributions of Ken Gray, Technical Project Manager III at Lumen, whose insights and expertise enriched the analysis and direction of this work.

The views expressed in this paper are those of the author and do not necessarily reflect those of the contributors or their organizations.

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.

Footnotes

¹ Multi-source data prepared by 4MC Partners for Lumen analysis.

² We also consider fiber availability (an investment pool to which we contribute) and constraints in the U.S. construction labor force.

³ The Register, *Goldman Sachs warns AI bubble could burst datacenter boom*, September 2, 2025.

⁴ CBRE, *North America Data Center Trends H2 2024*, February 2025.

⁵ Atlantic-ACM, *ACM Market Intelligence*, 2024.

⁶ RBC, *RBC ImagineTM: Multi-Industry Generative AI Update and Perspectives*, 2025.

Why Lumen?

For enterprises navigating the demands of Cloud 2.0, Lumen offers a strategic advantage that goes beyond connectivity. With one of the world’s most advanced fiber infrastructures and API-driven network solutions, Lumen empowers your enterprise to deploy, scale and secure mission-critical workloads with speed and confidence. Our unified management and direct access to leading cloud providers enable rapid innovation and seamless integration—so your organization can stay ahead of the AI-driven curve.