

Designing Enterprise AI Infrastructure: Lessons From Hyperscaler AI Networks

3 August 2026 - ID G00848028 - 14 min read

By: Nauman Raja, Andrew Lerner

Initiatives: Strategy, Risks and Opportunities; Reinvent I&O as an Enabler of AI Value; Architect the Next-Generation Enterprise For A Tech-Intensive Future; Delivery of Functional Responsibilities; Capitalize on Disruption With Emerging Technologies; Technologies and Markets

Hyperscalers are redefining AI data center networking. Analysis of the AI network fabrics used by major hyperscalers reveals the design principles and operational practices that enterprises can adopt to improve the performance, reliability and efficiency of their own AI infrastructure.

Insights at a Glance

Hyperscalers have redefined AI data center networking by treating the network as a system-level constraint rather than a supporting component. While enterprises cannot replicate hyperscale architectures, heads of I&O can adopt the design principles that will improve their own AI infrastructure's performance, reliability and efficiency.

Key Advice

Adopt hyperscaler networking practices selectively.

- Treat scale-up infrastructure as a vendor-integrated product
- Focus design effort on the scale-out network
- Prioritize workload-aware architectures, reliability and observability over peak performance metrics

Core Insights

AI infrastructure is constrained by networking, not compute: As accelerator performance has increased, network architecture has become the primary determinant of utilization, reliability and value.

Reliability matters more than peak performance: Hyperscalers optimize for goodput — useful work completed between failures. Fault isolation, congestion management and rapid recovery have greater impact on AI performance than incremental bandwidth improvements.

Workload characteristics should drive network design: Training, inference and mixture of experts (MoE) workloads generate different traffic patterns and require different network behaviors. A single architecture will not optimize every AI use case.

Strategic Planning Assumption

By the end of 2027, more than half of data center switching spend will support AI workloads.

Impact

AI infrastructure is no longer constrained by accelerator performance alone. As enterprises scale AI training and inference environments on-premises, the network increasingly determines utilization, reliability and overall investment efficiency.

- AI workloads generate traffic patterns that differ significantly from traditional enterprise applications, exposing limitations in existing data center network designs.
- Network bottlenecks can reduce accelerator utilization, increasing the cost of AI initiatives without improving business outcomes.
- Hyperscalers have spent the last decade solving these problems at scale, creating a body of architectural and operational practices that enterprises can selectively adopt.
- The shift toward reasoning models, long-context inference and **agentic AI** is increasing east-west traffic and making network design more consequential.
- Organizations that treat AI networking as a strategic design decision can improve infrastructure efficiency, reduce operational risk and delay unnecessary capital expenditure.

The opportunity for heads of I&O is not to replicate hyperscale environments, but to apply proven networking practices that improve performance, resilience and scalability while avoiding the cost and complexity of hyperscaler-specific architectures.

Actions

- **Deploy a dedicated AI back-end network.** Separate accelerator-to-accelerator traffic from management, storage and user-facing networks to improve performance predictability and simplify troubleshooting.
- **Standardize on Ethernet-based AI fabrics.** Use Ethernet with RDMA (RoCE) as the default architecture for AI scale-out networks and focus on congestion management, load balancing and operational simplicity.
- **Increase network sophistication as deployments expand.** Start with vendor reference architectures for small deployments and introduce advanced scale-out capabilities only when cluster size and workload demands justify them.
- **Invest in telemetry and observability.** Establish visibility into latency, congestion, packet loss and workload utilization to identify performance bottlenecks before scaling infrastructure.
- **Treat scale-up as a procurement decision and scale-out as a design decision.** Use validated vendor architectures for scale-up and focus engineering effort on scale-out topology, resiliency and operational processes.

Cautions

- **Do not reuse general-purpose data center networks for AI workloads.** AI traffic patterns are fundamentally different from traditional enterprise applications and can expose congestion and oversubscription issues that do not appear in conventional environments.
- **Do not assume one network architecture fits all AI use cases.** Training, fine-tuning and inference generate different traffic patterns and may require different capacity, topology and operational models.
- **Do not overengineer early-stage deployments.** Avoid introducing advanced topologies, multiplane fabrics or hyperscale operational practices before scale and workload requirements justify them.
- **Do not optimize for peak performance alone.** Bandwidth and latency specifications do not guarantee effective AI performance. Prioritize reliability, fault isolation and sustained utilization.
- **Do not attempt to invest in hyperscaler architectures.** Custom silicon, optical circuit switching, co-packaged optics and multisite training fabrics solve problems most enterprises do not have and add unnecessary cost and complexity.

How to Execute

AI infrastructure scaling has shifted from compute to networking. Once accelerators reached sufficient performance, the primary constraint became how effectively thousands of chips communicate. As a result, network design now determines system performance, reliability and efficiency.

Hyperscalers have already adapted. They treat the network as a first-class system component, optimize for useful work completed between failures (“goodput”), and design fabrics around workload behavior rather than peak bandwidth. Their architectures prioritize scale-out design, fault isolation and workload-aware traffic patterns.

While most enterprises lack the financial, technical and personnel resources of a hyperscaler by at least an order of magnitude, there are some architectural, procedural and technical best practices that enterprises **can and should** emulate. This research identifies the patterns observed within hyperscalers that are both relevant, impactful and doable for enterprises.

AI Infrastructure Has Shifted From Compute to Network Constraints

AI infrastructure no longer scales by adding more compute. It scales by improving how computing chips are connected. Once accelerators reached sufficient performance, the limiting factor became the network that coordinates them (see [Market Guide for AI Network Fabrics](#)).

The shift from compute-bound to network-bound systems changes how infrastructure must be designed and operated. Peak accelerator performance is no longer the primary measure of system capability. Effective performance depends on the network's ability to sustain communication under load, handle failures, and adapt to workload patterns.

Hyperscaler Landscape and Network Architectures

Hyperscalers operate AI infrastructure at a scale that forces different design decisions. While implementations vary, their approaches converge on a small number of consistent patterns.

First, they optimize for **goodput**, defined as useful work completed between failures. At large scale, failure is constant, and sustained performance matters more than peak throughput.

Second, they treat **scale-up as a tightly integrated system** delivered by the vendor, while expressing system-level behavior through the **scale-out network**. Topology, congestion control and load balancing determine how efficiently performance scales beyond a single rack.

Third, they design networks around **workload characteristics**. Training workloads generate synchronized, collective traffic, while inference workloads increasingly produce sustained, point-to-point flows. Network behavior is tuned to match these patterns.

Finally, they assume failure as a normal condition. Network designs prioritize fault isolation, rapid rerouting and maintaining progress under degraded conditions.

Three Shifts That Are Reshaping AI Data Center Networks

The design decisions made by hyperscalers reflect broader changes in how AI infrastructure operates. These shifts redefine the priorities for data center networking and determine which architectural choices remain viable at scale.

1. Ethernet Has Become the Default AI Fabric

AI networking has largely consolidated around Ethernet and RDMA (RoCE). While InfiniBand remains relevant for some frontier-scale NVIDIA clusters, most hyperscalers have standardized on Ethernet-based scale-out networks due to cost, ecosystem maturity and operational flexibility.

2. Reliability Determines Effective Performance

At AI scale, failures are expected. The most successful hyperscaler architectures focus less on peak throughput and more on maintaining useful work during hardware faults, congestion events and degraded conditions.

This has elevated **goodput** over theoretical performance as the primary measure of system efficiency.

3. Workload Characteristics Drive Network Design

Training, inference and mixture of experts (MoE) models generate fundamentally different traffic patterns. Network architectures optimized for one workload may perform poorly for another.

Hyperscalers increasingly design networks around workload communication patterns rather than assuming a single architecture can support all AI use cases equally well.

Implications for I&O: What to Adopt and What to Avoid

Hyperscaler designs solve problems at a scale most enterprises do not operate. Their value lies in the principles they apply, not the specific architectures they deploy. Heads of I&O should adopt the elements that transfer directly to enterprise environments and avoid those that introduce unnecessary complexity.

What to Adopt

These design patterns are consistently observed across hyperscaler environments and translate effectively to enterprise-scale deployments.

Standardize on Ethernet-Based Scale-Out

Ethernet with RDMA over converged Ethernet (RoCEv2) is the practical foundation for AI networking today. It provides sufficient performance, a mature ecosystem and operational familiarity. It also aligns with the direction of industry standards and vendor roadmaps (see [Focus Your Effort Where It Matters: The Scale-Out Layer](#)).

Separate the AI Back-End Network

The accelerator-to-accelerator fabric should operate as a dedicated network. Training and inference traffic should not share bandwidth with front-end, storage or management traffic.

This enables:

- Predictable performance under load
- Isolation of congestion and faults
- Simplified troubleshooting and tuning

Use Rail-Optimized or Structured Topologies

AI workloads benefit from predictable, structured connectivity between nodes. Designs that maintain alignment between accelerators and network paths reduce interference and improve collective performance.

This approach scales effectively from small clusters to multirack deployments.

Prioritize Congestion Control and Load Balancing

AI traffic patterns stress traditional network behaviors. Large, sustained flows and synchronized communication require:

- Network design aligned to workload characteristics
- Adaptive routing or multipath load balancing
- Effective congestion control mechanisms
- Careful buffer and queue management

These factors have a greater impact on performance than incremental increases in bandwidth.

Invest in Telemetry and Observability

AI workloads amplify the impact of small faults. Limited visibility makes performance issues difficult to diagnose and resolve.

Effective environments require:

- End-to-end visibility across compute and network
- Monitoring of congestion, latency and packet loss
- Tooling to correlate network behavior with workload performance

Design Around Workload Use Cases

Network requirements vary by workload:

- Collective-heavy training workloads require high east-west bandwidth
- Inference workloads require efficient point-to-point communication and sustained throughput

Capacity planning and topology decisions should reflect the dominant workload pattern rather than theoretical peak demand.

What to Avoid

These patterns reflect hyperscaler-specific solutions or emerging technologies that do not translate effectively to enterprise environments.

Avoid Reusing Existing Data Center Networks for AI Workloads

Traditional enterprise networks were not designed for AI traffic patterns. AI workloads generate large volumes of east-west traffic, synchronized communications and sustained flows that can overwhelm fabrics built for conventional applications.

Deploy dedicated AI networking infrastructure for accelerator-to-accelerator communication rather than sharing production data center networks.

Custom Silicon and Proprietary Interconnects

Hyperscalers invest in custom interconnects and silicon to control cost and performance at extreme scale. Enterprises should rely on vendor-integrated systems for scale-up and avoid attempting to replicate these designs.

Advanced Optical Architectures

Technologies such as optical circuit switching, co-packaged optics or specialized fiber are designed for hyperscale or experimental deployments.

They introduce:

- Operational complexity
- Limited vendor support
- Uncertain cost-benefit at enterprise scale

For most environments, pluggable optics and short-reach copper remain sufficient.

Ultralarge Nonblocking Fabrics

Fully nonblocking fabrics across very large clusters are expensive and unnecessary in most enterprise environments.

Some level of oversubscription is acceptable and often desirable, as AI workloads do not uniformly consume all available bandwidth.

Cross-Data Center Training Fabrics

Extending a single training job across multiple sites introduces latency, coordination complexity and operational risk.

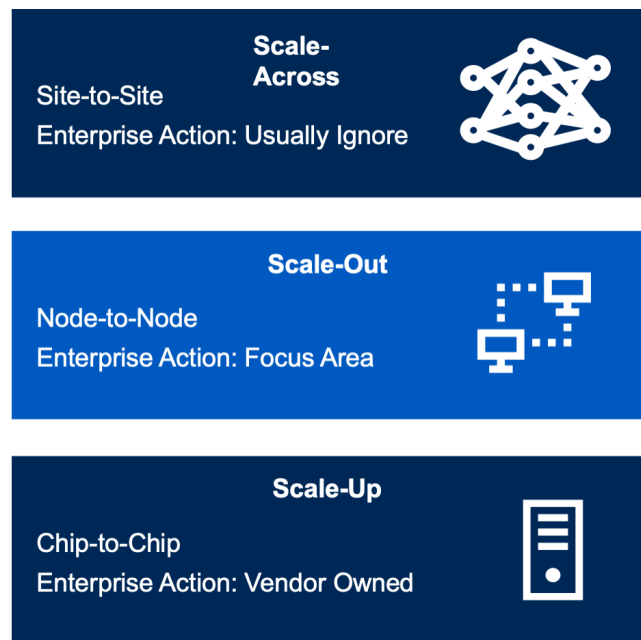
Most enterprises should treat AI deployments as single-site systems for training and use standard multiregion patterns for inference distribution.

Focus Your Effort Where It Matters: The Scale-Out Layer

AI infrastructure has three network layers: scale-up, scale-out and scale-across. Of these, only scale-out represents a meaningful design surface for most I&O teams. It is where performance, cost and reliability are determined once deployments extend beyond a single rack. Figure 1 below shows the network layers found in AI and where to focus as an enterprise.

Figure 1: Hyperscaler AI Network Fabric (AINF) Layers

Hyperscaler AINF



Source: Gartner
848028

Gartner

Scale-Up Is a Purchased System

Scale-up infrastructure is delivered as an integrated product:

- GPU, interconnect and switching are predefined
- Performance characteristics are fixed by the vendor
- Integration and compatibility are already validated

This layer should be treated as a consumption decision, not a design problem. Attempting to modify or optimize scale-up provides limited benefit and increases operational risk.

Scale-Out Determines System Performance

Once workloads extend beyond a single scale-up domain:

- Network design defines effective utilization
- Congestion and imbalance become primary constraints
- Additional compute capacity does not guarantee higher performance

Scale-out becomes the dominant factor in:

- Training efficiency
- Inference throughput
- Overall system stability

Decisions made at this layer directly affect whether deployed accelerators operate at expected efficiency.

Key Design Decisions at Scale-Out

Effective scale-out design requires deliberate choices in a small set of areas:

- Topology: structured, predictable designs that align network paths with workload behavior
- Load balancing: mechanisms that distribute large flows effectively across available paths
- Congestion control: maintaining stability under sustained and synchronized traffic
- Oversubscription: balancing cost and performance rather than defaulting to nonblocking designs
- Fault isolation: limiting the impact of failures on the wider system

These decisions are interdependent and must be considered as a system.

Scale-Out Scales With the Environment

Scale-out design requirements change with cluster size:

- Within a rack: minimal or no scale-out design required
- Across a few racks: structured topology and basic tuning become important
- Across many racks: full design discipline is required, including resilience and observability

This progression allows organizations to increase design complexity only as needed, rather than overengineering from the outset.

Scale-Across Is Not Required for Most Environments

Connecting multiple data centers into a single AI cluster is a hyperscale requirement.

For most enterprises:

- Training workloads should remain within a single site
- Inference distribution should follow standard multiregion architectures

Design effort should not be spent on cross-site fabrics unless operating at extreme scale.

I&O teams do not need to design the entire system. They need to focus on the one layer where their decisions matter. That layer is scale-out.

Final Recommendations

Heads of I&O should focus on a small number of deliberate design and operational decisions. These recommendations prioritize system efficiency, reliability and practical deployment outcomes.

1. Default to Ethernet-Based Scale-Out

Use Ethernet with RoCE as the standard for AI cluster networking. It provides sufficient performance and a broad ecosystem, and aligns with vendor and industry direction.

2. Treat the Back-end AI Fabric as a Dedicated Network

Separate accelerator-to-accelerator traffic from all other network functions. Do not share the fabric with front-end, storage or management traffic.

3. Prioritize Reliability and Stability Over Peak Performance

Design for sustained performance under load. Focus on:

- Fault isolation
- Traffic stability
- Consistent throughput

Evaluate infrastructure based on effective utilization, not theoretical peak.

4. Treat Scale-Up as a Procurement Decision, Not a Design Problem

Select validated vendor systems for scale-up. Avoid attempting to optimize or redesign this layer.

Concentrate design effort on scale-out.

5. Make Explicit Design Choices in the Scale-Out Network

Make explicit decisions on:

- Topology
- Load balancing
- Congestion control
- Oversubscription

Do not rely on default configurations.

Representative Vendors in AI Network Fabrics

See Table 1 below for a list of sample vendors in the AI network fabric (AINF) market. This list has been sourced from the [Market Guide for AI Network Fabrics](#).

Table 1: AINF Vendors

(Enlarged table in Appendix)

Vendor	Platform or product brand
Aria Networks	Deep Networking
Arista Networks	Etherlink
Arrcus	ACE-AI
Astrape Networks	Eefiniti
Aviz Networks	ONES
Cisco	Hyperfabric, Nexus
Dell	PowerSwitch
DriveNets	AI Networking Fabric
Edge-Core Networks	AI Switch
Eridu	Eridu AI network switch
H3C	H3C AI Fabric
Hedgehog	Hedgehog Fabric
HPE (Juniper)	PTX, QFX
Huawei	Xinghe
Netris	Controller
Nexthop	AI Networking
Nokia	Data Center Fabric
NVIDIA	SpectrumX, Quantum
Ruijie	AI Fabric
Unifabrix	Max
Upscale	Skyhammer

Source: Gartner (August 2026)

Success Measures

AI Infrastructure Goodput

Measure: Percentage of AI workload execution time spent performing useful work rather than recovering from failures, waiting for resources or experiencing infrastructure bottlenecks.

Target: More than 90% for production AI environments.

Why it matters: Goodput is the most direct measure of AI infrastructure effectiveness. It captures the combined impact of network reliability, congestion management, observability and operational processes.

AI Infrastructure Observability Coverage

Measure: Percentage of AI infrastructure components (compute, network and storage) covered by end-to-end telemetry and monitoring.

Target: One hundred percent visibility across the AI infrastructure stack.

Why it matters: Goodput cannot be improved if organizations cannot identify packet loss, congestion, hardware failures or workload bottlenecks. Observability is a prerequisite for optimization.

Network-Related AI Performance Incidents

Measure: Percentage of AI workload slowdowns, failures or restarts attributable to network congestion, packet loss or fabric instability.

Target: Reduce network-related performance incidents by 50% within 12 months.

Why it matters: Hyperscalers focus heavily on fault isolation and resilience because small network failures can have outsized impacts on AI workload performance.

Workload-Aligned AI Network Architectures

Measure: Percentage of AI environments with documented network designs aligned to workload requirements (training, inference or mixed-mode).

Target: One hundred percent of new AI deployments include workload-specific network design and capacity planning.

Why it matters: Training, inference and MoE workloads generate different traffic patterns. Aligning network architecture to workload requirements improves utilization, scalability and cost efficiency.

Success Statement

Success is achieved when organizations can accurately observe AI infrastructure performance, rapidly identify and isolate network issues, and sustain high levels of AI workload goodput.

Evidence

The following hyperscalers were examined: AWS, Google, Microsoft, Meta, Alibaba, ByteDance and Oracle Cloud Infrastructure (OCI). These hyperscalers vary in their level of vertical integration, use of proprietary technology and workload focus.

Recommended by the Authors

Some documents may not be available as part of your current Gartner subscription.

[Market Guide for AI Network Fabrics](#)

© 2026 Gartner, Inc. and/or its affiliates. All rights reserved. Gartner is a registered trademark of Gartner, Inc. and its affiliates. This publication may not be reproduced or distributed in any form without Gartner's prior written permission. It consists of the opinions of Gartner's Business and Technology Insights Organization, which should not be construed as statements of fact. While the information contained in this publication has been obtained from sources believed to be reliable, Gartner disclaims all warranties as to the accuracy, completeness or adequacy of such information. Although Gartner insights may address legal and financial issues, Gartner does not provide legal or investment advice and its insights should not be construed or used as such. Your access and use of this publication are governed by [Gartner's Usage Policy](#). Gartner prides itself on its reputation for independence and objectivity. Its insights is produced independently by its Business and Technology Insights Organization without input or influence from any third party. For further information, see "[Guiding Principles on Independence and Objectivity](#)." Gartner insights may not be used as input into or for the training or development of generative artificial intelligence, machine learning, algorithms, software, or related technologies.

Table 1: AINF Vendors

Vendor	Platform or product brand
Aria Networks	Deep Networking
Arista Networks	Etherlink
Arrcus	ACE-AI
Astrape Networks	Eefiniti
Aviz Networks	ONES
Cisco	Hyperfabric, Nexus
Dell	PowerSwitch
DriveNets	AI Networking Fabric
Edge-Core Networks	AI Switch
Eridu	Eridu AI network switch
H3C	H3C AI Fabric
Hedgehog	Hedgehog Fabric
HPE (Juniper)	PTX, QFX
Huawei	Xinghe
Netris	Controller

Nexthop	AI Networking
Nokia	Data Center Fabric
NVIDIA	SpectrumX, Quantum
Ruijie	AI Fabric
Unifabrix	Max
Upscale	Skyhammer

Source: Gartner (August 2026)