

Ethernet for the Era of Agentic AI

OCP Educational Webinar Series
September 25, 2025



Community-driven hyperscale innovation for all

Ethernet for the Era of Agentic AI

New Requirements, High-Speed Networking Innovation & Testing Strategies

James Kelly

VP Market Intelligence
and Innovation
OCP



Sameh Boujelbene

VP, Ethernet Switch—Data Center and AI
Networks for AI Workloads Market Research
Dell'Oro Group



Aniket Khosla

VP, Wireline Product
Management
Spirent



Agenda Overview

Evolving Requirements of AI Networks

Ethernet is Winning in AI Networks

Testing Challenges in AI Networks

Innovative Testing Strategies

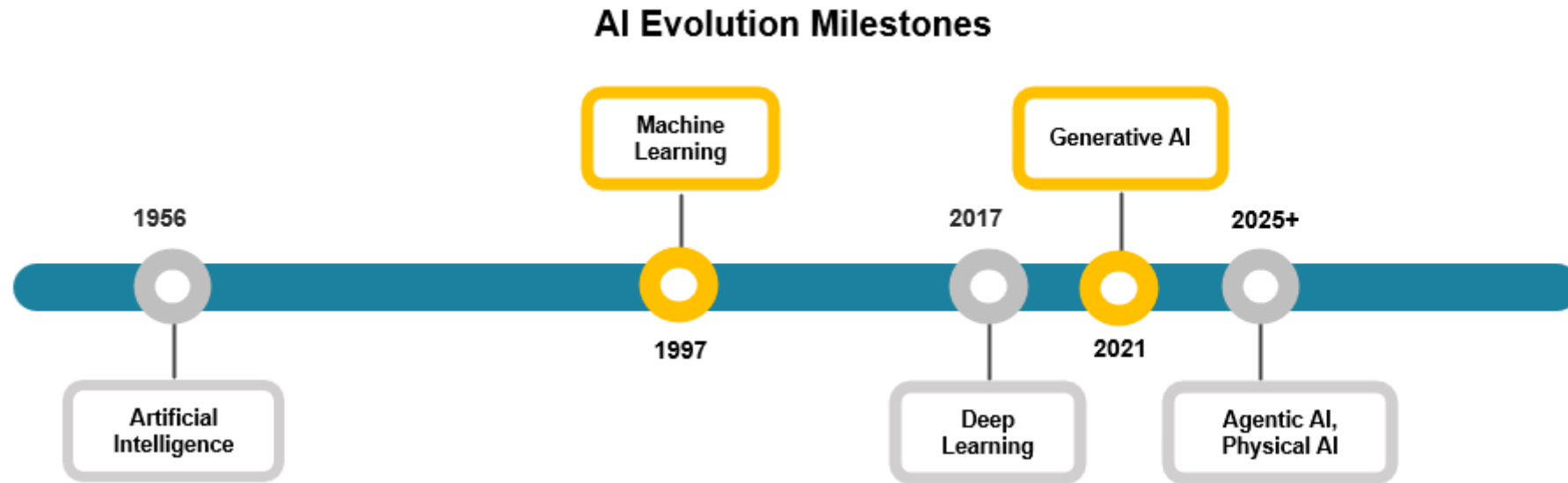
Advancing Ethernet for AI and HPC

Summary

Live Q&A

AI Keeps Changing and Evolving at a Rapid Pace

- New age of AI reasoning driving new scaling laws: from pre-training to post-training and test-time scaling
- Jensen Huang at GTC25: “Demand for compute is 100X what we predicted a year ago”

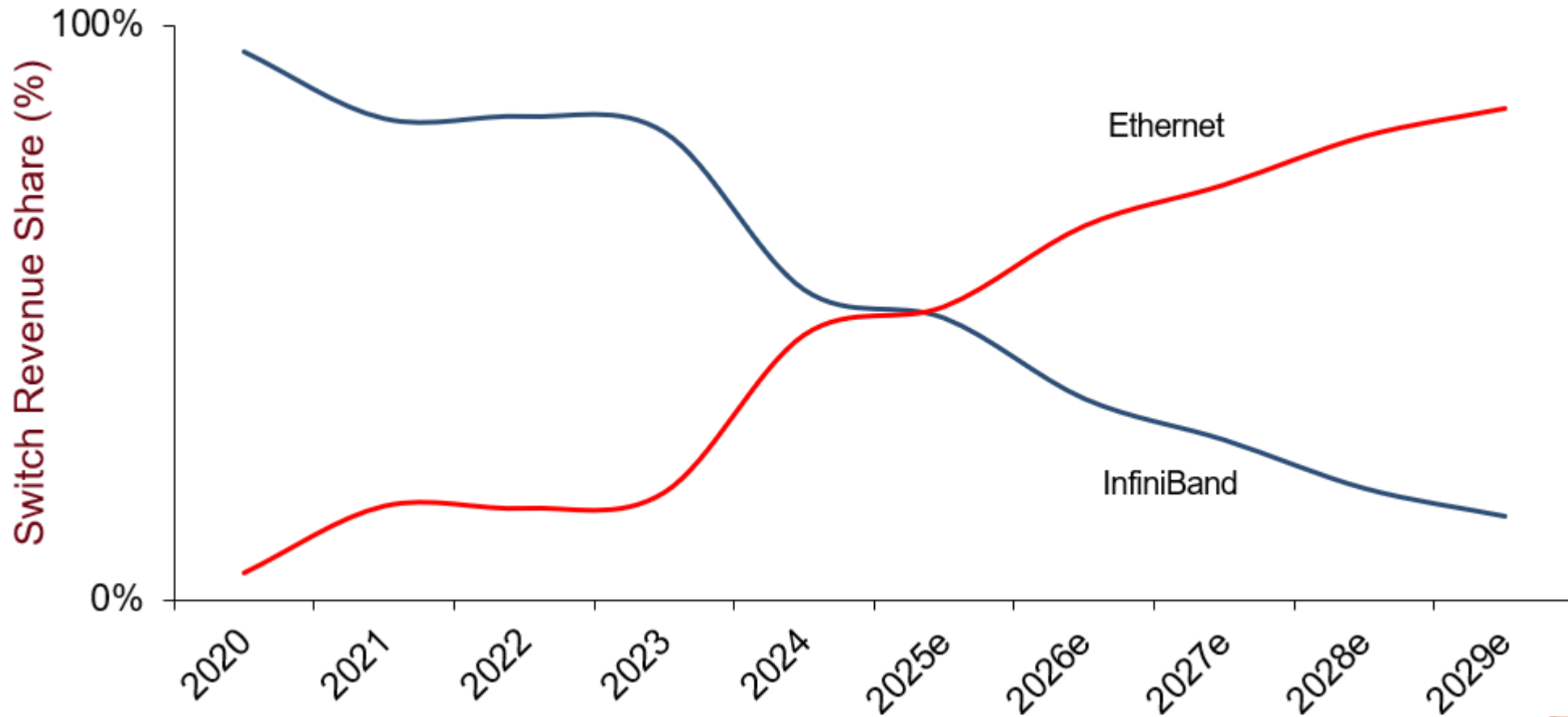


AI Clusters Exploding in Size: the Network is the Computer

	Pre-2023	2023	2024	2026+
Cluster Size (XPU's)	1 K	25 K	100 K	1 M
	↓ 2X	↓ 3X	↓ 5X	↓ 10X
Interconnects	2K	75 K	500 K	10 M

- Rising AI cluster scale brings new requirements:
 - Scalable, open and flexible network fabric to support rapid growth and evolving architectures
 - Vendor diversity to foster faster innovation and risk mitigation
 - Cost efficiency and broad availability to ensure sustainable scaling of AI infrastructure

A Very Fast Migration... from InfiniBand to Ethernet in AI Back-End Networks

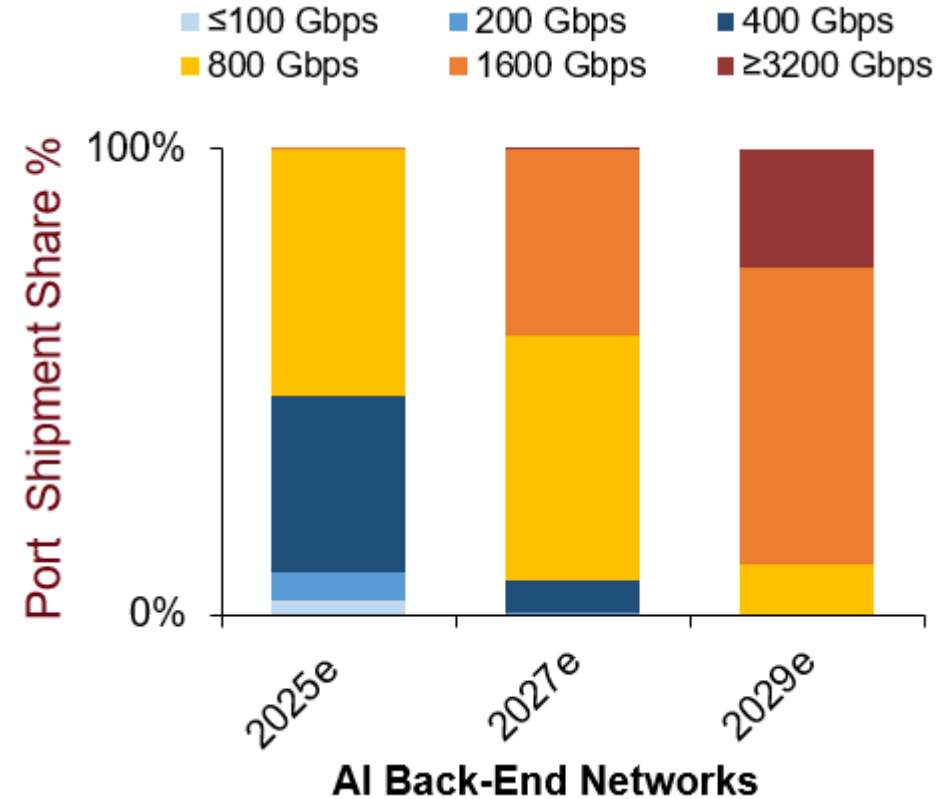
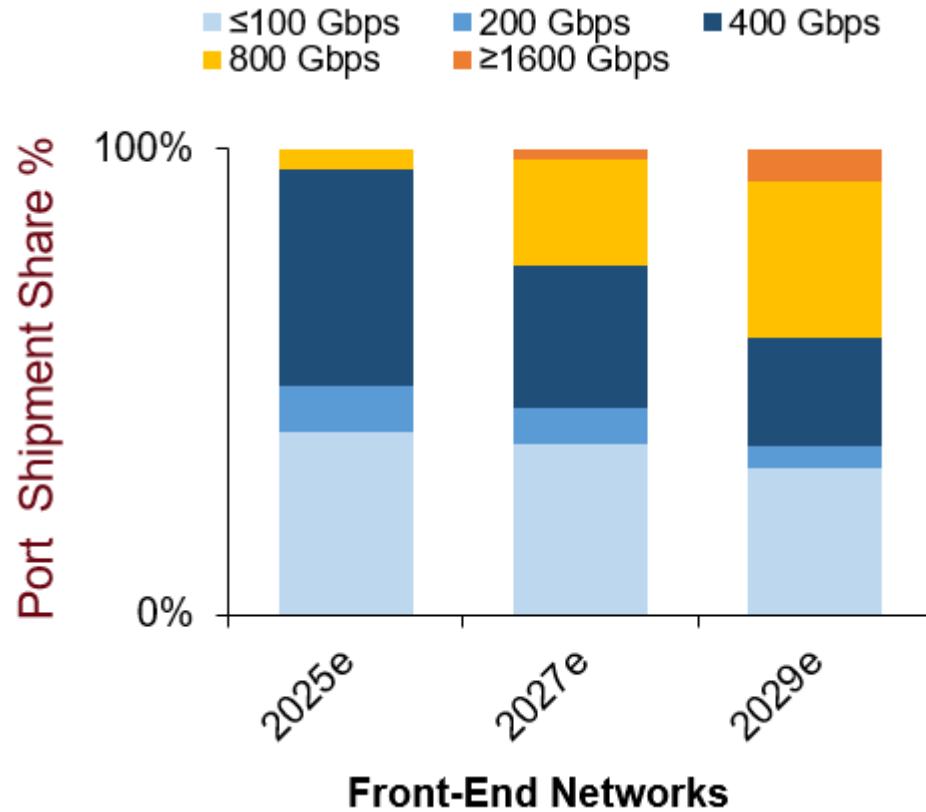


Source: Dell'Oro AI Networks for AI Workload Report
Scale-out only, excludes scale-up



Speed Transition in AI Clusters

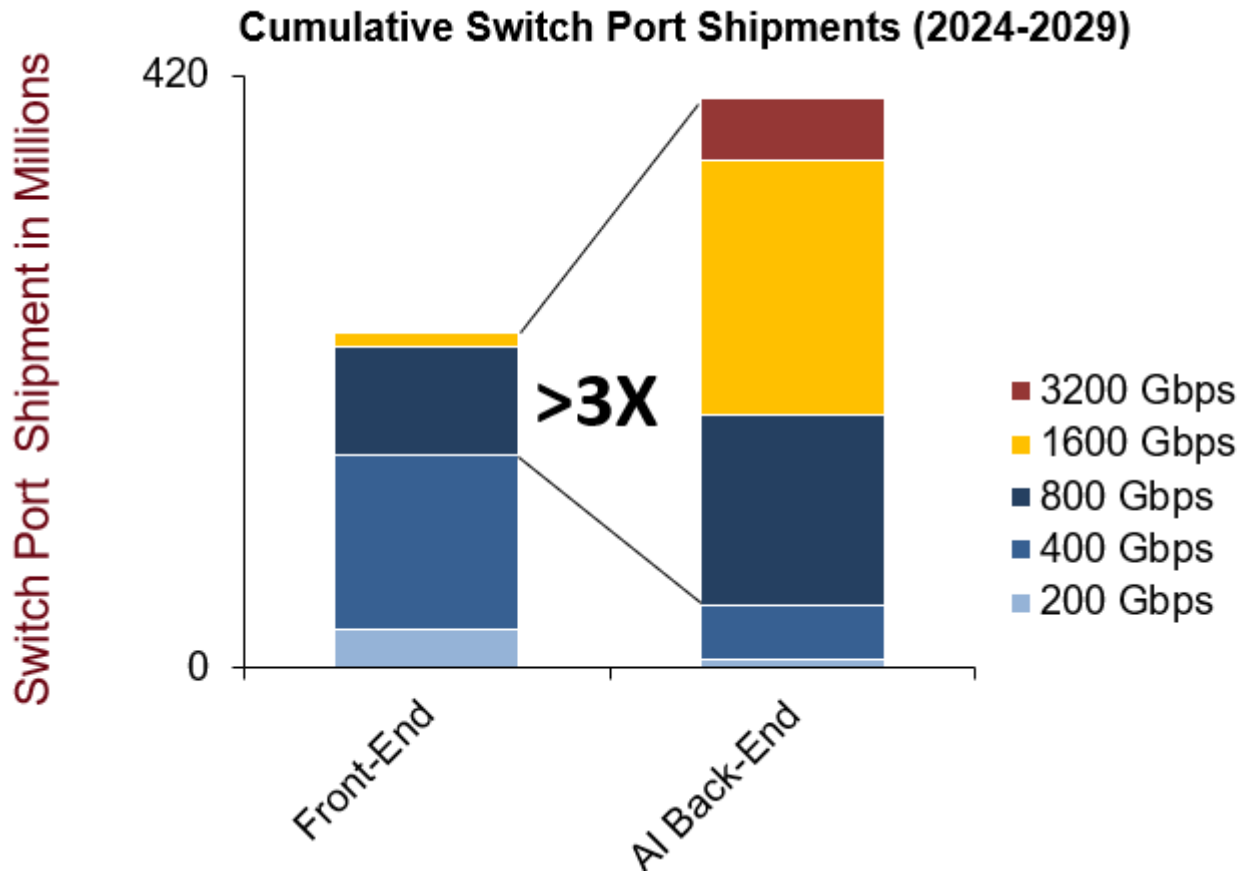
Twice as Fast as in the Front-end Network



Source: Dell'Oro AI Networks for AI Workload Report
Scale-out only, excludes scale-up



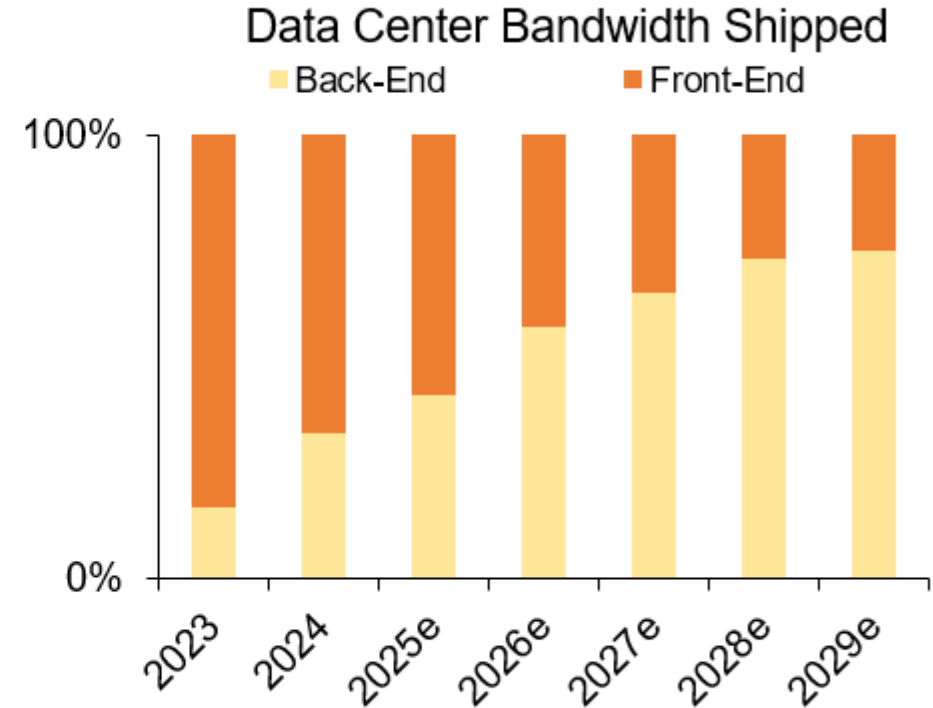
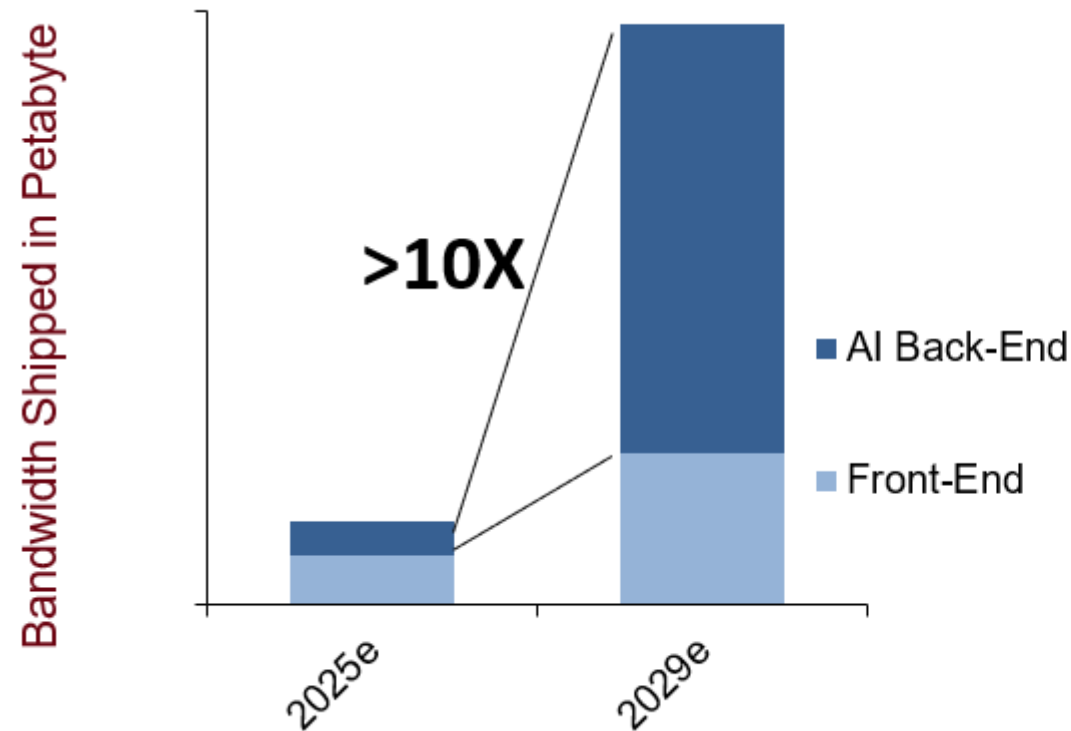
AI Clusters Triple the Demand for High-Speed Ports in Front-End Networks



- 800/1600/3200 Gbps shipments (2024-2029):
 - ✓ Back-end: well over 300 M
 - ✓ Front-end: under 100 M

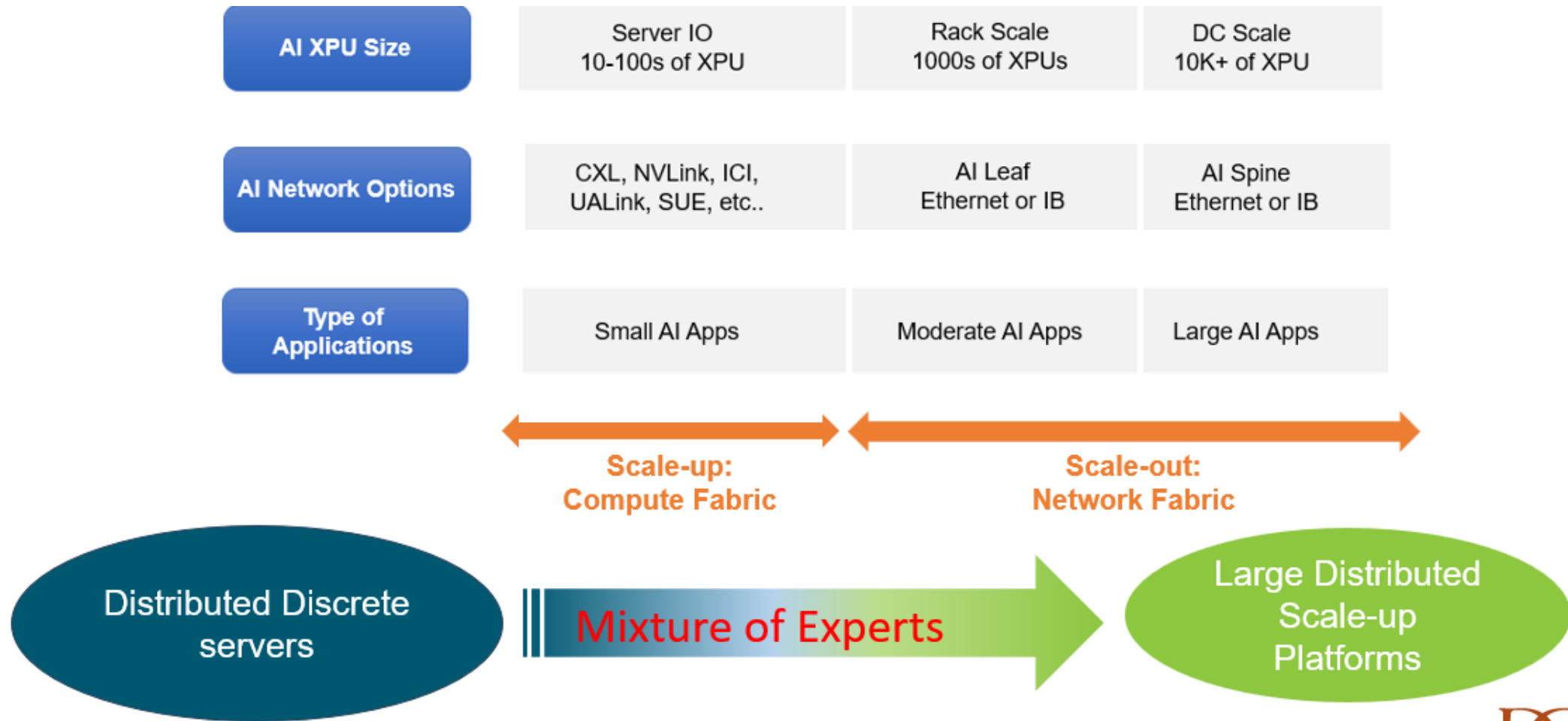
Source: Dell'Oro AI Networks for AI Workload Report
Scale-out only, excludes scale-up

Bandwidth Requirement in AI Clusters



Source: Dell'Oro AI Networks for AI Workload Report
Scale-out only, excludes scale-up

Scale-up vs. Scale-out Network Design Options for AI Clusters



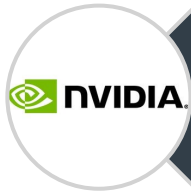
XPU could be GPU, TPU or any other type of accelerators

Poll Question 1

Q: Which statement most closely matches your organization's position regarding Ethernet and InfiniBand connectivity for advanced AI traffic?

Responses
A. We have been using InfiniBand and don't envision moving away from it in the next 4-5 years.
B. We have been using InfiniBand but planning to move to Ethernet in the next 4-5 years.
C. We are already using some combination of both high-speed Ethernet and InfiniBand.
D. We are not sure how to proceed, but we favor high-speed Ethernet .
E. We are not sure how to proceed, but we favor InfiniBand .
F. We are currently using mostly public cloud providers for our AI infrastructure.

Major Innovations Under Way



Unveils Groundbreaking Blackwell GPUs

Delivering an astounding 20 petaflops of AI performance.
Boasts a remarkable 25-fold improvement in energy efficiency.



AI-as-a-Service European Expansion

Commits \$2.2 Billion to meet surging demand.
Bringing NVIDIA Blackwell to the region for the first time at scale.



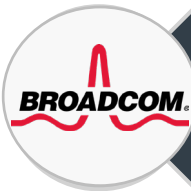
Trainium and Inferentia

Amazon's purpose built Trainium chips optimized for training deep learning models. Inferentia to generate inferences more quickly.



AI as a Service - Expands Service Offerings

Powers One-Click-Simple global deployment for AI applications
AI is now easier and cost effective for customers to deploy.



Announces Tomahawk Ultra Ethernet Switch

Delivers ultra-low latency and 77B packets/sec throughput.
Optimized for AI/HPC scale-up environments.



Announces Partnership

Integration of NVIDIA's AI capabilities with Cisco's hardware.
Creation of advanced AI-powered networking solutions.



Introduces "Apple Intelligence" & M4 Processor

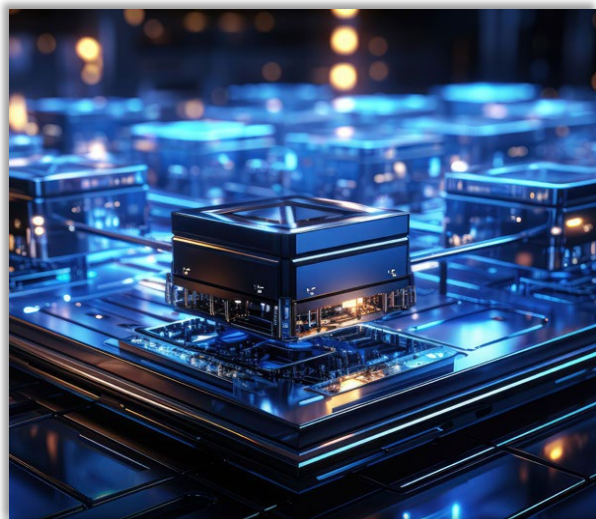
AI Integrated across iOS 18, iPadOS 18, and macOS Sequoia.
M4 is Apple's fastest neural engine ever.



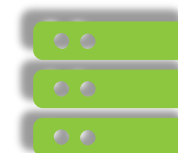
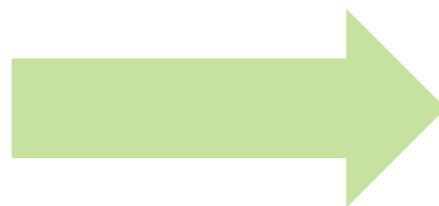
Announces new Etherlink AI Switching Platform

Optimized network performance for large AI clusters.
Utilizes Broadcom Tomahawk 5 silicon, capacity of 51.2Tbps.

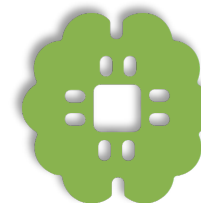
AI Applications and Infrastructure



Model Complexity & Size



+ Compute



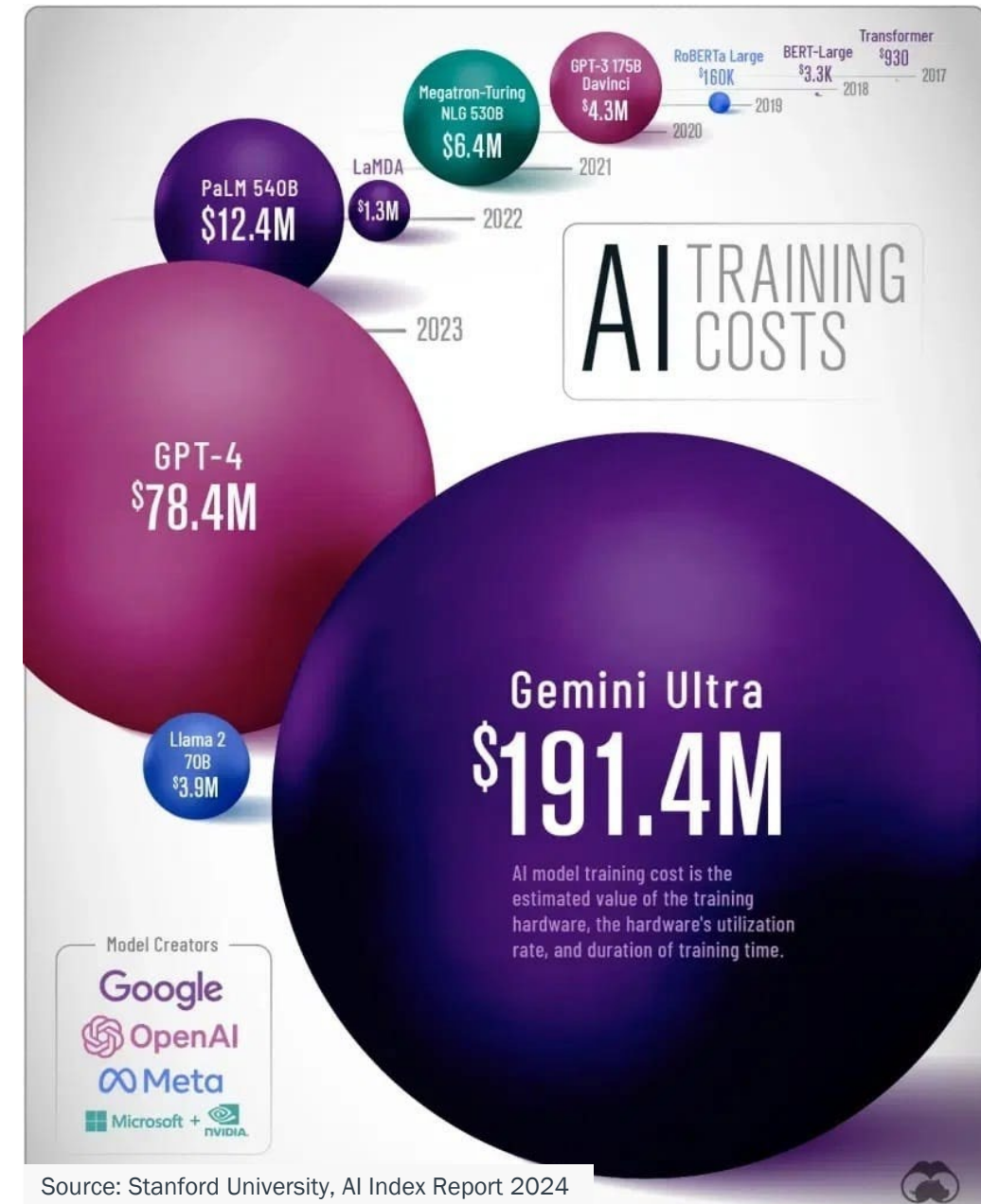
+ Memory



+ Network

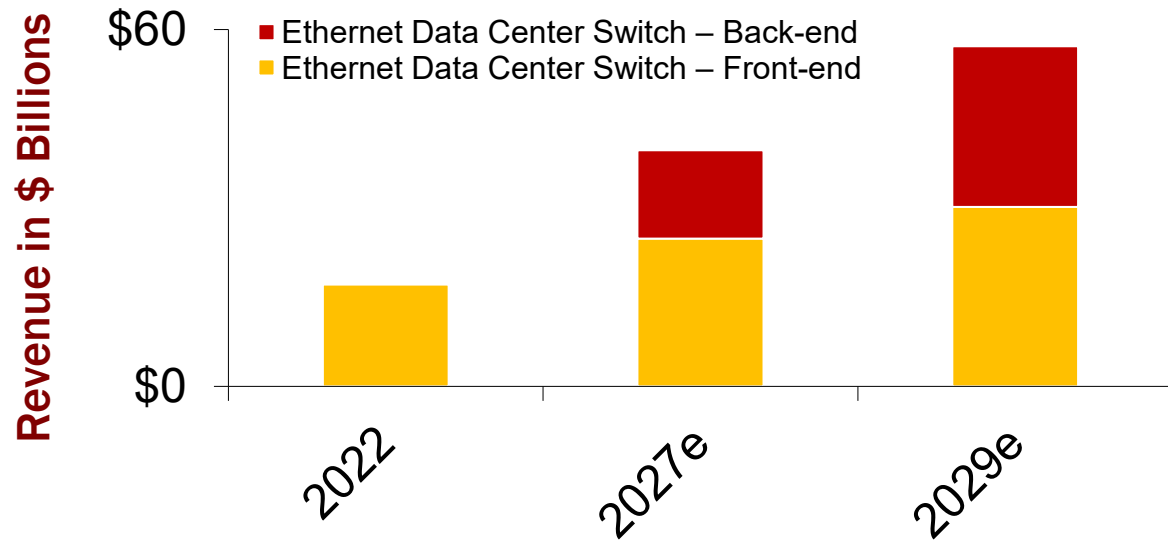
Cost of AI

- AI Training costs for frontier models were increasing at an **annual rate of 2–3x**
 - The largest models could exceed \$1 billion by 2027 if trends continue
- “DeepSeek Moment” has decreased the expected per-unit cost of AI
- AI innovations are vastly more accessible now, expanding the overall market & accelerating adoption
- Total industry **spend is expected to grow** even as per-unit costs fall

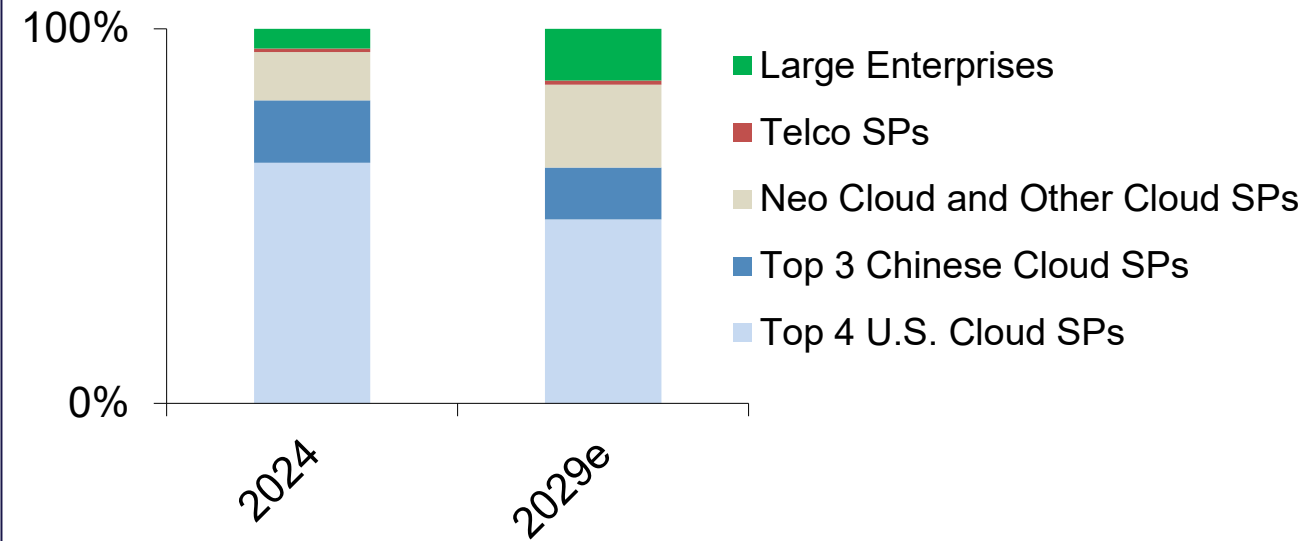


AI Data Center Networks Market Segments & Size

AI Back-end Networks Doubling the TAM for Ethernet Data Center Switching



Growing Contribution from Neo Cloud SPs and Enterprise Customers



Source: Dell'Oro Group

Networking Requirements: Training & Inference in AI Workflows

Every stage of AI workflows has specific networking requirements

AI Model Training

Large-scale, distributed training requirements

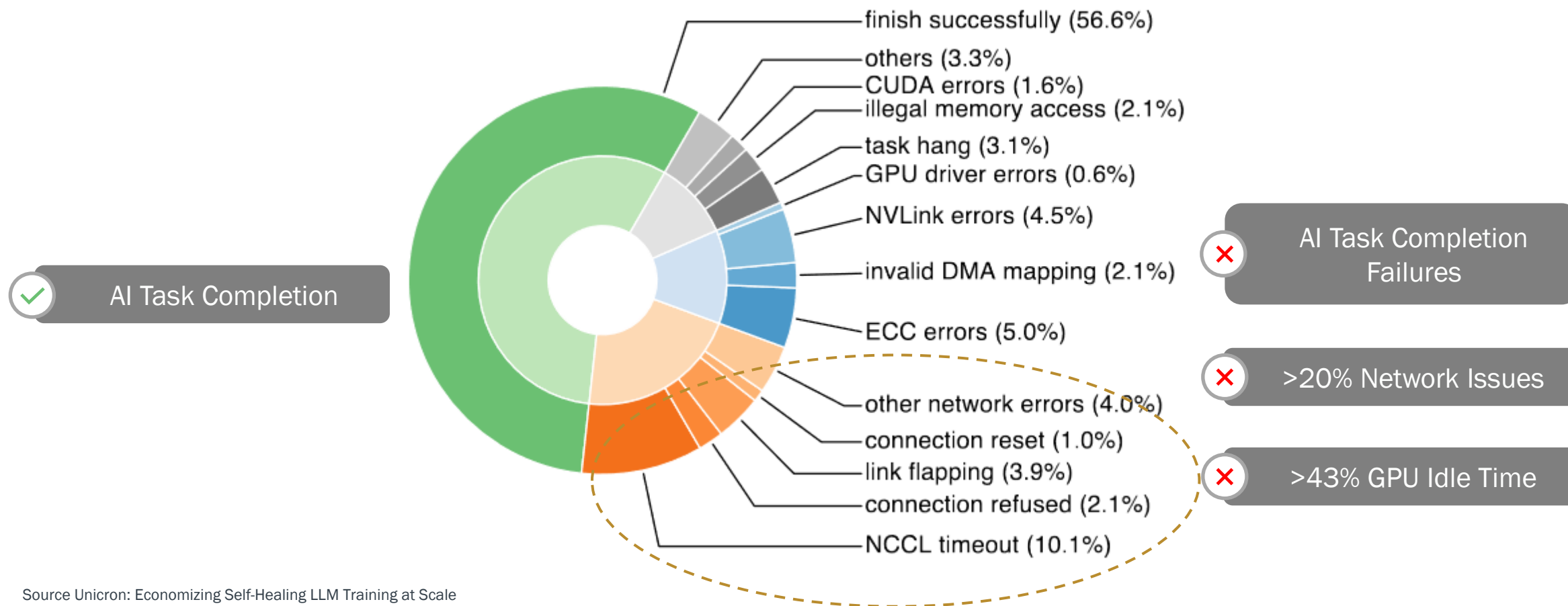
- Efficient communication, data transfer, and synchronization between nodes
- High bandwidth
 - 400G, 800G Ethernet or InfiniBand networks
- Deterministic latency
 - Minimal latency jitter
- Lossless connectivity
 - No packet loss

AI Model Inference

Focusing on efficient data flow, low latency & serving a potentially high volume of requests

- Ultra-low latency communication - near real-time financial fraud, chatbot use cases
- Load balancing to handle multiple simultaneous requests
- Elastically scaling compute infrastructure
- Robust security and compliance across data & operations
- High-performance connectivity far-edge, mid & regional data center sites - closer to the end user

The Network Impact on AI Data Centers



Source Unicorn: Economizing Self-Healing LLM Training at Scale
Tao He, Xue Li, Zhibin Wang, Kun Qian, Jingbo Xu, Wenyuan Yu, Jingren Zhou
Alibaba Group & Nanjing University, December 2023

Impact of Low GPU Utilization

Increasing complexity by 10x increases the need for GPUs by 70x (not sustainable)

Improve GPU Utilization = Save Big + Improve AI Ops

Model name	Model Size (billion params)	Dataset size (billion tokens)	Training ZFLOPS (10 ²¹)	GPU	GPU FLOPS	GPU utilization	Training time (in weeks)	Number of GPUs
OPT	175	300	430	H100	3,000	0.5	1	474
<u>LLaMA</u>	65	1,400	600	H100	3,000	0.5	1	662
<u>LLaMA 2</u>	34	2,000	400	H100	3,000	0.5	1	441
LLaMA2	70	2,000	800	H100	3,000	0.5	1	882
GPT-3	175	300	420	H100	3,000	0.5	1	463
GPT-4* (est.)	1,500	2,600	31,200	H100	3,000	0.5	1	34,392

Table 3 — GPU cluster size for training various LLM models with H100 GPUs, FP8 arithmetic, and 50% utilization.

Memory Access

Parallelism Optimization

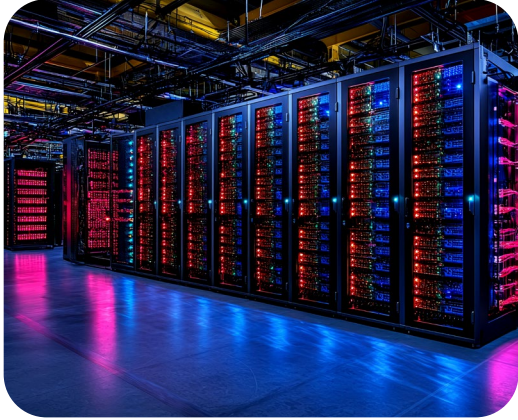
Network Bottlenecks

< 1% Packet Loss =
33% Loss in AI Training Performance

>70x more GPUs needed
due to low utilization

Source: GPU Fabrics for GenAI Workloads, [Sharada Yeluri](#) Dec 2023

Validating AI Networking Performance



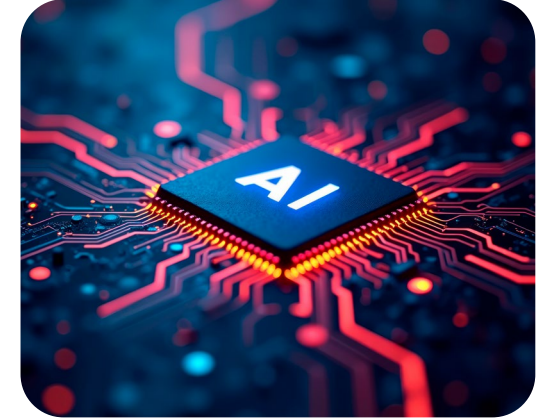
Back-End AI Network

- Going from 400Gbps today to 1.6Tbps Ethernet per GPU by 2027
- New Ethernet features to support AI demanding workloads



Front-End AI Network

- Connect the GPU nodes to the data center network for inferencing at 400Gbps and 800Gbps speeds
- Shared network with Storage and Applications



Agentic Workflows

- Inference is driving network upgrades in back-end & front-end networks to support AI-driven customer experiences
- AI-based apps and agents will need a scalable infrastructure to deliver low latency & a high volume of requests securely

Preparing Networks for Optimal AI Data Center ROI

Testing AI Fabric: Challenges With Real xPU Testing

High Cost & Limited Repeatable Testing

- Procurement
- Deployment
- Management
- Test Creation/Execution
- Power Consumption
- Results/Analytics and Ongoing Test Development



It takes a data center to test a data center



Burden is on end customer at all levels of the solution



Knowledge of AI network landscape, traffic & protocols required

Poll Question 2

Q: Are you testing your AI data center prior to deployment? How much time do you spend on testing?

Responses
A. We are testing our AI data centers prior to deployment, testing takes 2-4 weeks .
B. We are testing our AI data centers prior to deployment, testing takes a few days .
C. We are testing our AI data centers prior to deployment, testing takes a few hours .
D. We are not testing our AI data centers prior to deployment.
E. We are not deploying an AI data center .

Testing Approaches for Validating AI Infrastructure



Expertise

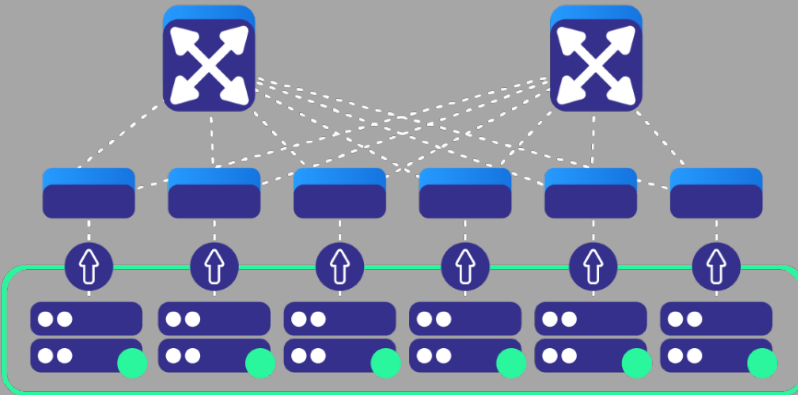


Accuracy & Performance



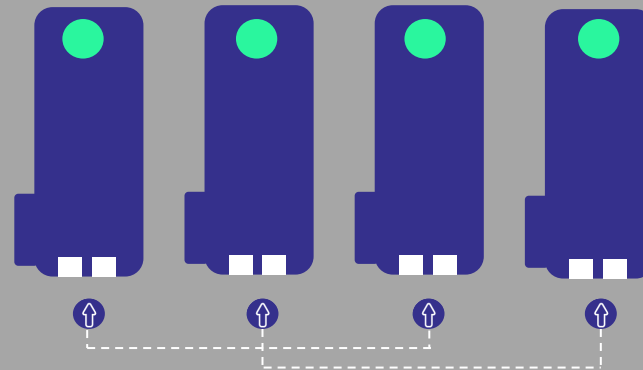
Real-World Emulation

xPU Emulation



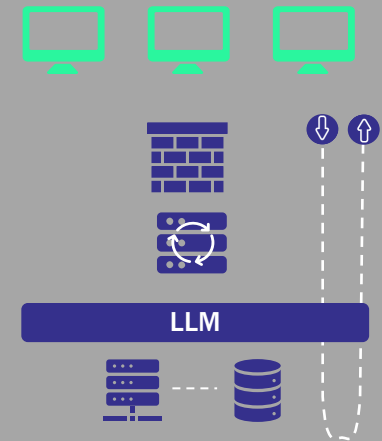
AI Data Center
RoCEv2 Network Fabric

RDMA Agents



AI Data Center
RoCE NIC / DPU

Client/User Emulation



AI Datacenter LLM
Inference Infrastructure

Case Study

US/EU Neocloud provider validates front-end and 800G back-end AI networking

CHALLENGE

- Neocloud provider needed to ensure client GPU cluster access thru high performance front-end Ethernet & IP routing network.
- Needed to prepare for addition from 400G InfiniBand back-end only to 800G Ethernet to support higher scale and potential additional AI offerings.

SOLUTION

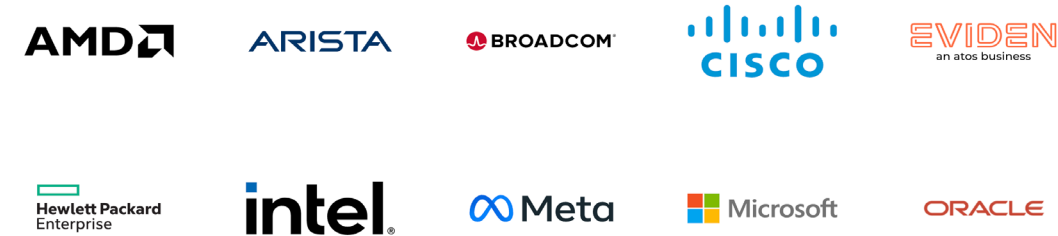
- Provided new B3 800G solution with both Ethernet/Routing and RoCEv2/CCL AI Networking capabilities to validate front-end and back-end networks.

IMPACT

- Customer was able to **validate optimal performance** of front-end network, providing reliable multi-client access to GPU clusters and resources.
- **Sped up deployment of 800G Ethernet** and was able to demonstrate to clients it performs just as well as current InfiniBand Back-end.



Steering Members



General Members



The Emergence of UEC

Ultra Ethernet Consortium (UEC) Mission:

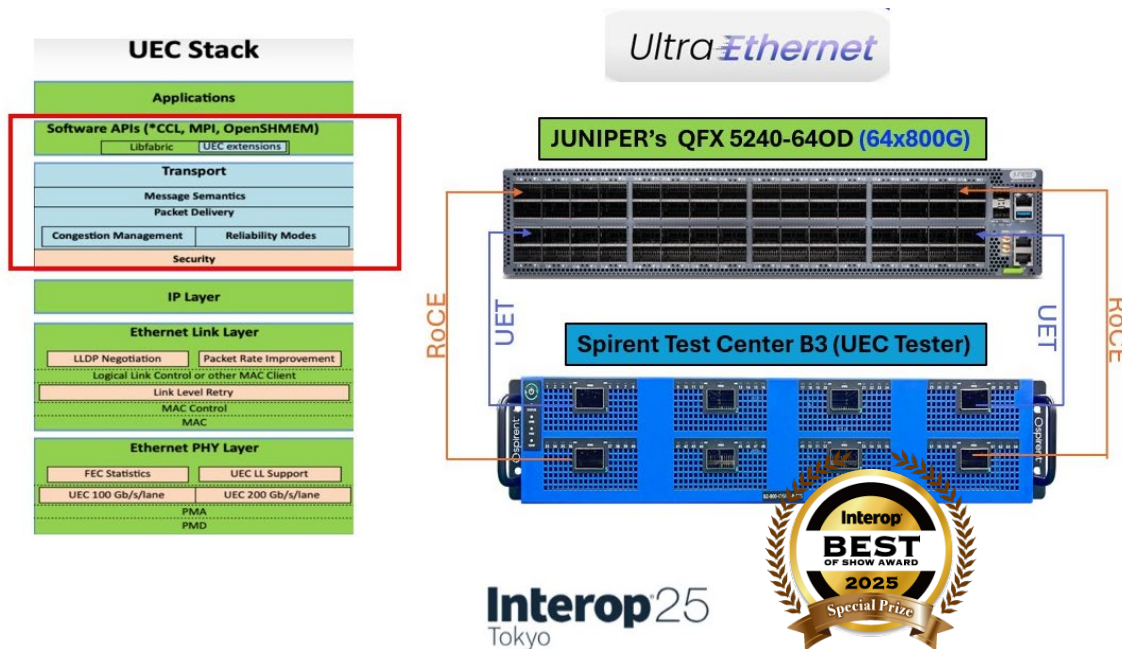
Deliver an Ethernet-based open, interoperable, high-performance, full-communications stack architecture to meet the growing network demands of AI & HPC at scale

We are a leading contributor to the UEC, helping shape an open, high-performance Ethernet architecture built to meet the massive scale and complexity of AI and HPC networks.

- Engaging early with UEC proposals to stay ahead of evolving standards
- Collaborating on specifications to enable first-to-market Ultra Ethernet test solutions
- Delivering to market UEC based assessment and validation solutions

World's First Ultra Ethernet Transport Validation

- At Interop Tokyo 2025, Spirent and Juniper Networks showcased the first public UET benchmark, using Juniper QFX5240-640D switch and Spirent B3 800G Appliance to co-run UET and RoCEv2 traffic in a realistic data center topology.
- This innovative setup earned a prestigious *ShowNet Special Prize*, marking a major milestone for the UEC ecosystem & validating UET with Juniper's software stack as production-ready for scalable AI & HPC fabrics.

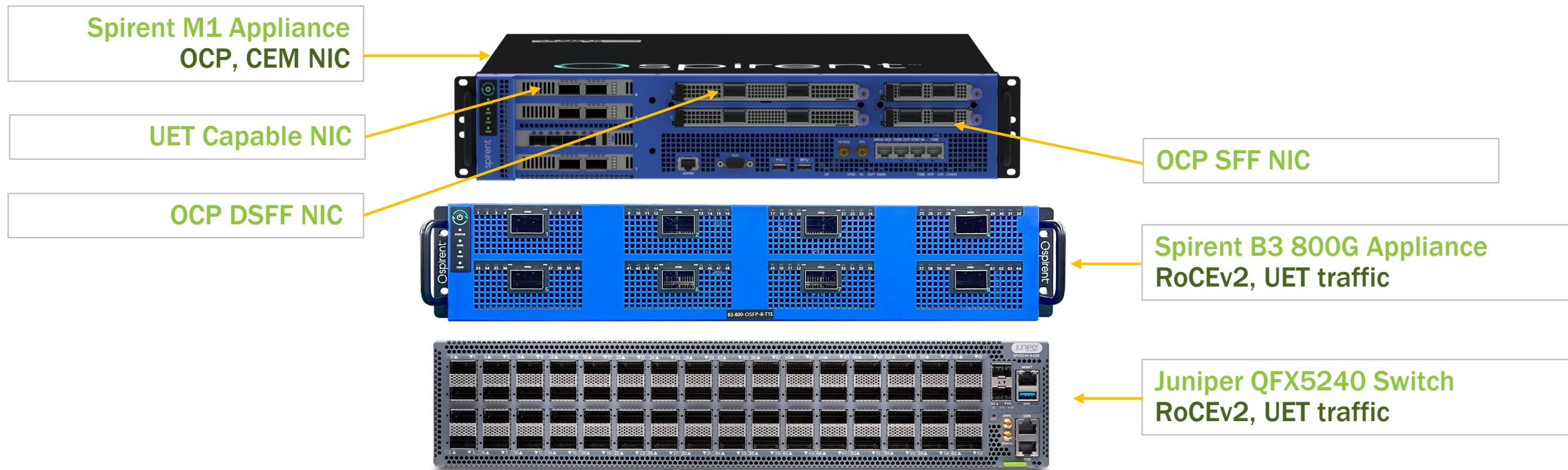


Key objectives included:

- Validate **UET packet forwarding** with real-time traffic
- Verify **key KPIs for AI data center workloads**
- Demonstrate **UET & ROCEv2 co-existence**

OCP Innovation Village Showcase

- Showcasing new packet-trimming capabilities and interoperability with emerging transport standards like UET, to ensure next-gen infrastructure can meet the evolving demands of AI-driven networking.
- Highlighting collaboration among OCP members to evaluate the performance of UEC-based products and underscore the need for advanced testing tools to address the demands of AI-scale networking.



Conclusion

- AI keeps changing and evolving at a very rapid pace
- Total AI industry spend is expected to continue to grow even as per unit costs continue to fall
- Rising AI cluster size demands scalable, open, and cost-efficient network infrastructure with vendor diversity to enable rapid growth, innovation, and sustainable scaling
- Ethernet is winning in AI back-end networks
- Speed migration in AI back-end networks is twice as fast as in the front-end networks
- Innovative testing strategies are key to accelerate AI deployments

Live Q&A

James Kelly

VP Market Intelligence
and Innovation
OCP



Sameh Boujelbene

VP, Ethernet Switch—Data Center and AI
Networks for AI Workloads Market Research
Dell'Oro Group



Aniket Khosla

VP, Wireline Product
Management
Spirent



Thank you for attending

Ethernet for the Era of Agentic AI

OCP Educational Webinar Series

September 25, 2025



Community-driven hyperscale innovation for all