



AIDC地端化與智慧校園

以RoCE v2 實現高速安全的資料交換

Jimmy Wang

Sep. 2025



願亡者得以安息
願傷者盡速痊癒
願災區重獲平安



AIDC地端化與智慧校園

以RoCE v2 實現高速安全的資料交換

Jimmy Wang

Sep. 2025



Climate of AI DC

AI 資料中心的網路需求

AI 大語言模型啟動第四次工業革命



Hugging Face

Models

1,484,462



Chatbot Arena
<https://lmarena.ai>

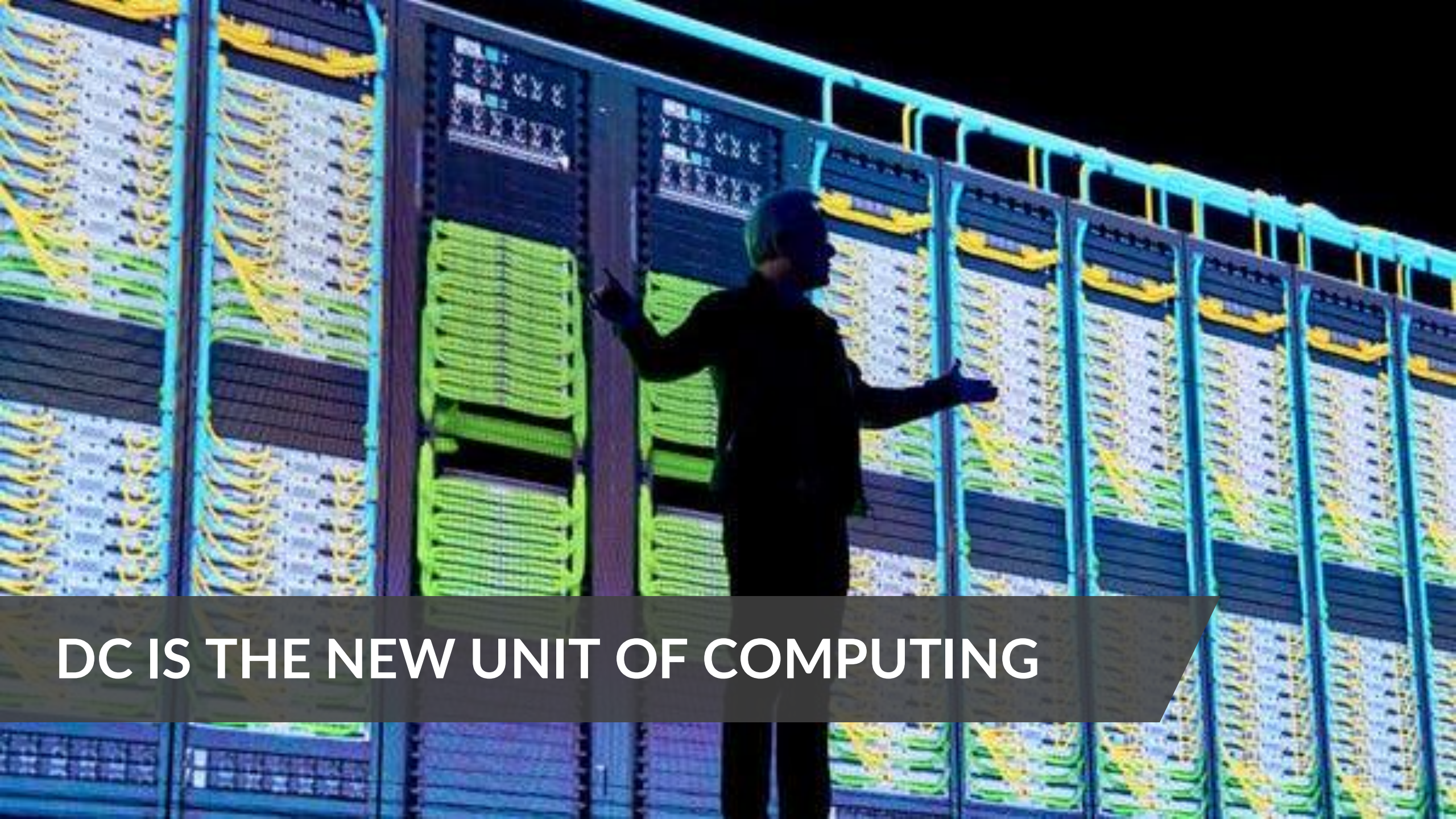


繁中LLM 聊天機器人競技場
<https://arena.twllm.com>



Rank* (UB)	Rank (StyleCtrl)	Model	Arena Score	95% CI	Votes	Organization	License	Knowledge Cutoff
1	1	chocolate (Early Grok-3)	1402	+7/-6	7829	xAI	Proprietary	Unknown
2	4	Gemini-2.0-Flash-Thinking-Exp-01-21	1385	+5/-5	13336	Google	Proprietary	Unknown
2	2	Gemini-2.0-Pro-Exp-02-05	1379	+5/-6	11197	Google	Proprietary	Unknown
2	1	ChatGPT-4o-latest (2025-01-29)	1377	+5/-6	10529	OpenAI	Proprietary	Unknown
5	2	DeepSeek-R1	1361	+8/-7	5079	DeepSeek	MIT	Unknown
5	8	Gemini-2.0-Flash-001	1356	+6/-5	9092	Google	Proprietary	Unknown
5	2	o1-2024-12-17	1353	+6/-5	15437	OpenAI	Proprietary	Unknown
8	6	o1-preview	1335	+4/-4	33169	OpenAI	Proprietary	2023/10
8	8	Qwen2.5-Max	1332	+7/-7	7370	Alibaba	Proprietary	Unknown
10	9	DeepSeek-V3	1317	+4/-4	17717	DeepSeek	DeepSeek	Unknown

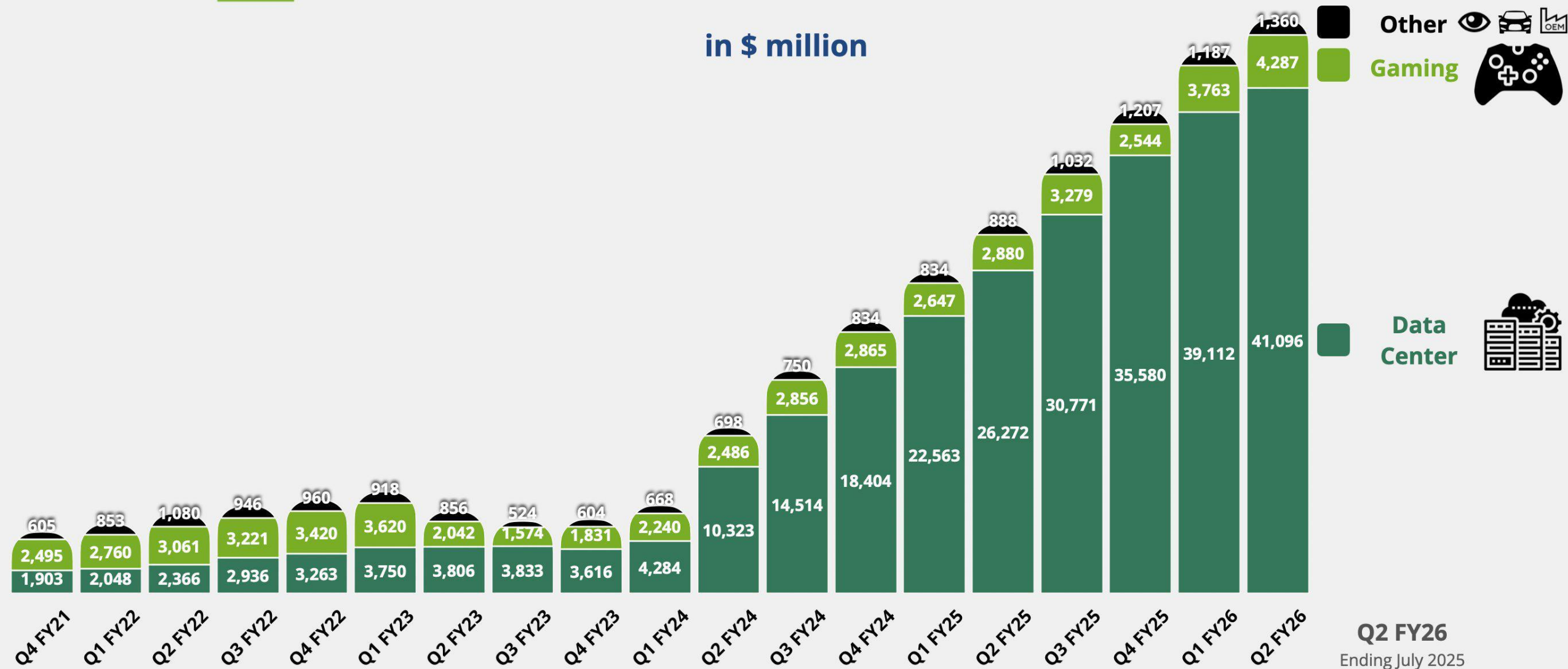




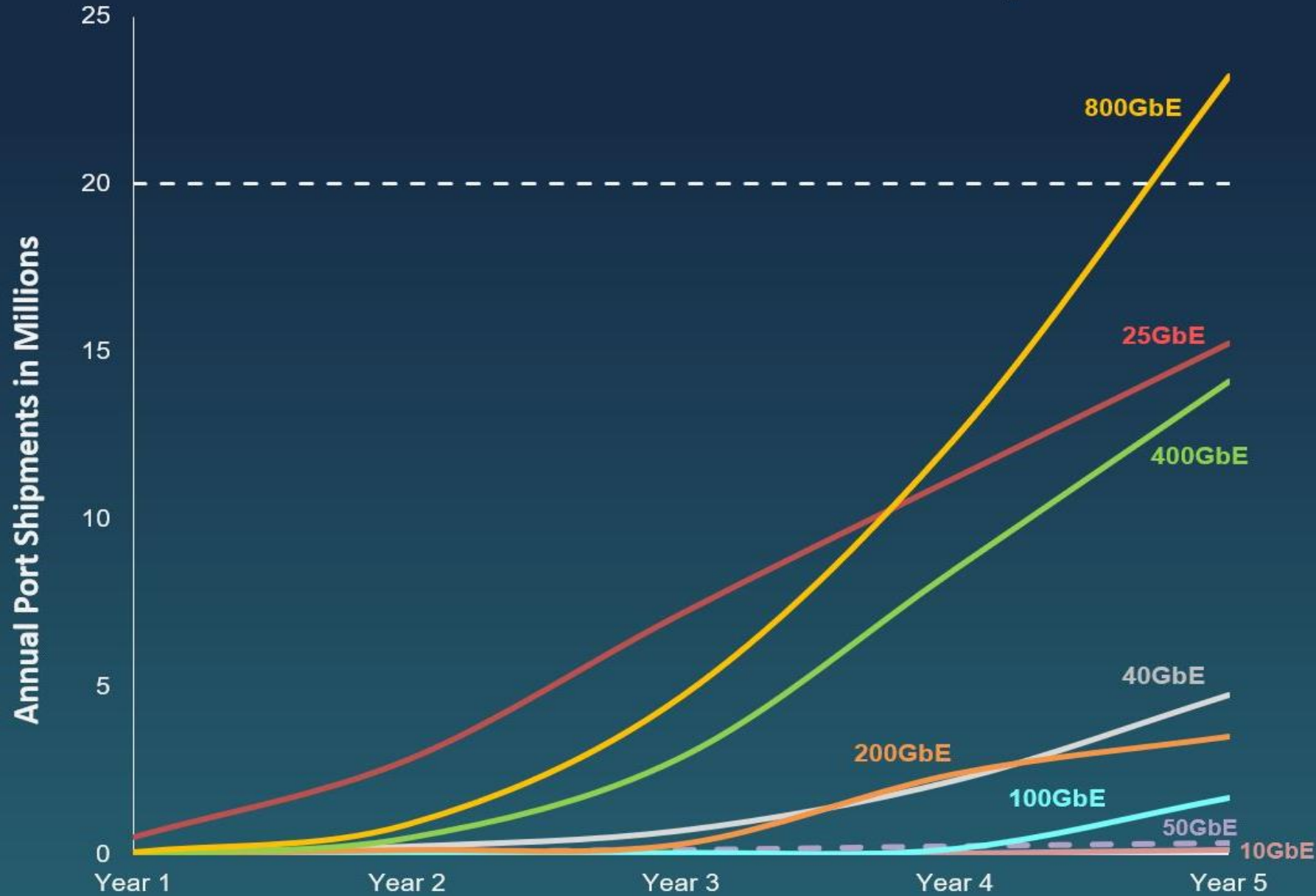
DC IS THE NEW UNIT OF COMPUTING

NVIDIA Revenue Breakdown

in \$ million



Data Center Ethernet Switch Annual Shipments



Year 1 Equals:

2023 for 800GbE

2016 for 25GbE

2018 for 400GbE

2011 for 40GbE

2019 for 200GbE

2012 for 100GbE

2020 for 50GbE

2001 for 10GbE



DATA CENTER NETWORKING MARKET

FORTUNE
BUSINESS INSIGHTS

Data Center Networking Market to grow at **13.1% CAGR** during 2025-2032



INDUSTRY DEVELOPMENT

Equinix partners with Singapore's GIC and Canada's CPPIB to expand data center infrastructure across the U.S.



NORTH AMERICA

Year	Market Size (Million)
2023	\$11.07
2024	\$11.95

Asia Pacific | Europe
Middle East & Africa | South America

TRENDS

Adoption of Multi-Cloud & Hybrid Cloud Networking in Enterprises

DRIVERS

Proliferation of Cloud Computing in Businesses



GLOBAL, BY INDUSTRY

IT & Telecom 21.3%
Media & Entertainment
Government & Defense
Manufacturing | BFSI
Healthcare | Retail
Others



BY COMPONENTS

Hardware | Software
Services

AI Data Center: Market view

InfiniBand is fading away, Industry is pivoting to Ethernet

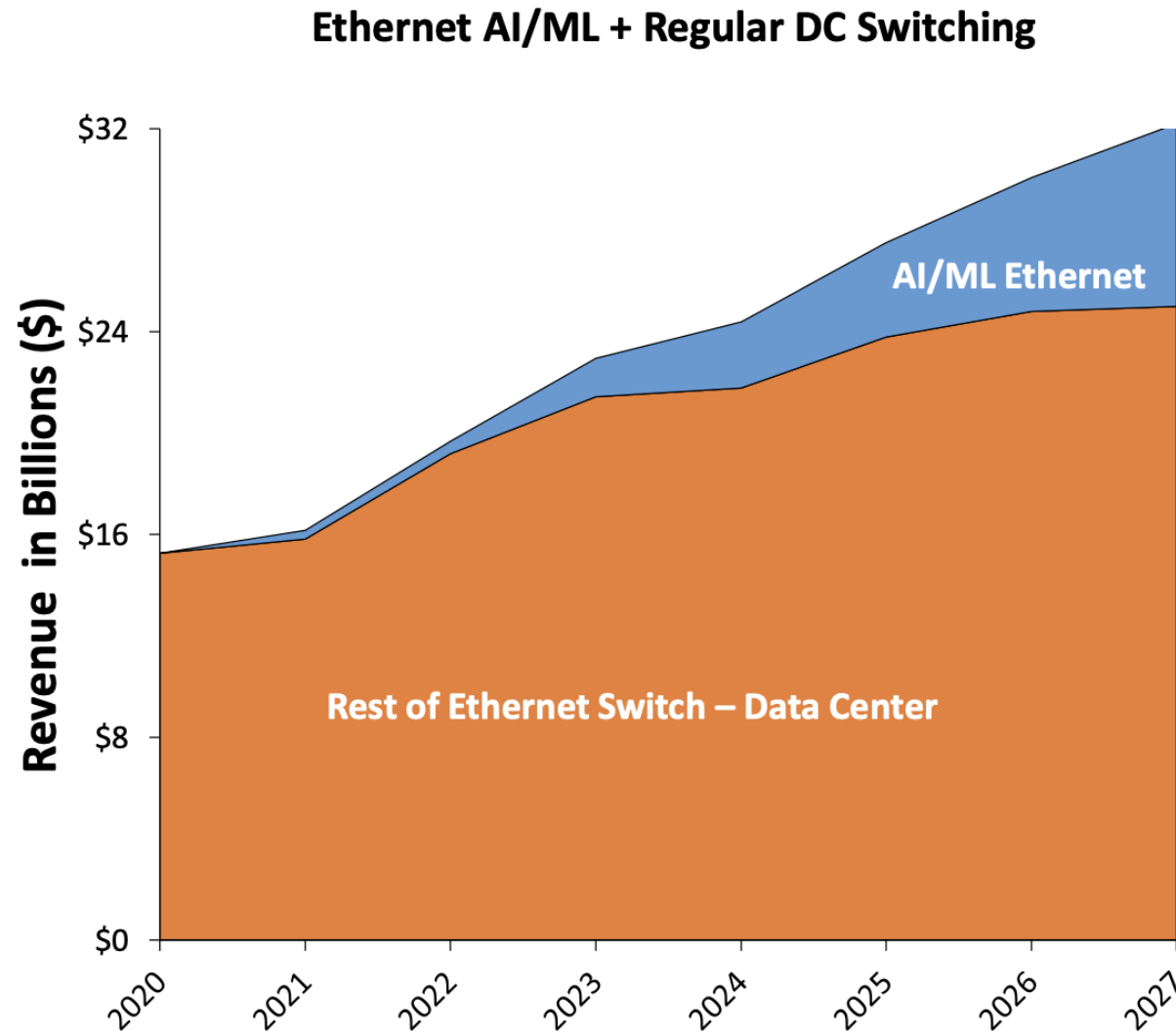
	2024	2025	2026	2027	2028
Share Revenue (%)					
Ethernet - Frontend	23%	30%	34%	30%	34%
Ethernet - Backend	23%	41%	44%	51%	50%
InfiniBand (Switching Only)	54%	29%	22%	19%	17%



DC Network CapEx: Market Transition with AI

Investments in some enterprise DCs are being paused...

to make way for investment in AI DCs



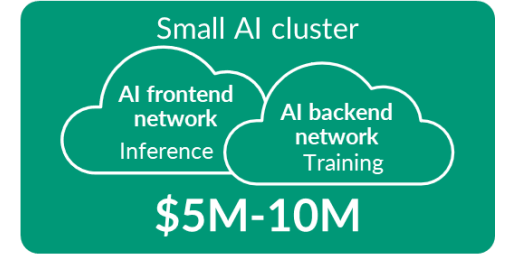
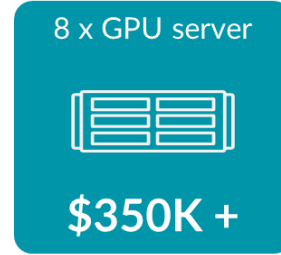
Source: 650 Group, 2023

Avoid an Infrastructure Bottleneck

If the network is a bottleneck that delays training job completion, expensive GPU time is wasted, and training becomes network-bound instead of compute-bound

Juniper Mission for AI Data Centers

Unlocking the full potential of AI with unparalleled network performance and ease of operations



GPUs represent the bulk of the CapEx investment



The network connects GPUs in distributed training

xAI

HIGH-PERFORMANCE BACKEND AI DATA CENTER NETWORKING THE CHALLENGE

- Accommodate large-scale and efficient GPU cluster and server connectivity.
- AI training is compute-intensive, with traffic flows that break traditional data center networking
- Efficient use of GPU cycles. (Job completion time)

THE SOLUTION

- Juniper QFX5240 for 800G leaf spine connectivity
- Enabling the same form factor for both Front and Back-end AI clusters
- Efficient Load balancing and congestion control for RDMA traffic

WHY JUNIPER?

- **Junos feature richness** for congestion management with ROCEv2
- **AI Lab for POC** for various load balancing scenarios
- **Lightening responsiveness to needs**
- **Better supply chain** for switch and optics with Juniper



Rami Rahim @ramirahim · Jul 24

Congrats @elonmusk @xAI @X! Excited for @JuniperNetworks to be a part of the Memphis Supercluster team and to bring our networking solutions to this innovative work.



Elon Musk @elonmusk · Jul 22

Nice work by @xAI team, @X team, @Nvidia & supporting companies getting Memphis Supercluster training started at ~4:20am local time.

With 100k liquid-cooled H100s on a single RDMA fabric, it's the most powerful AI training cluster in the world!

4

29

139

16K



AI Technology on Juniper

CORPORATE PROFILE

- Founded 2017; Series D (backed by SoftBank Google Ventures, BlackRock, et al.)
- Builds AI hardware and integrated systems to run AI applications from the data center to the cloud
- Purpose-built enterprise-scale AI platform is the technology backbone for the next generation of AI computing

THE SOLUTION

- Juniper QFX and PTX series fabric for high density, performance & scalability to move massive volumes of data
- Automation across design, deployment & operations of the network lifecycle with Juniper Apstra

THE RESULTS

- **Accelerate high-performance ML model building across industries** enabling customers to deploy AI in days, not months
- **Eases the construction of ML infrastructure & delivers enhanced capabilities and efficiency** for ML model training, inference, and high-performance computing
- **Five times better performance** than traditional GPU architecture

"In AI, data flow is king. We need the lowest network latency and the highest bandwidth possible, and the performance of the Juniper QFX5200 switches has been phenomenal,"

Vijay Tatkar, Director of Product Management

AI MANAGED SERVICES INFRASTRUCTURE



HIGH-PERFORMANCE AI INFRASTRUCTURE SOLUTIONS FOR ENTERPRISES

THE CHALLENGE

- Unfamiliarity with NVIDIA InfiniBand
- Lengthy lead times from Cisco, Arista and NVIDIA
- Pricing from Arista/Cisco was high and slow to respond to ionstream's questions

THE SOLUTION

- AI GPU pod features Juniper's cutting-edge QFX 5240 powering 800GbE infrastructure
- Data center networking foundation empowers flexible GPU-as-a-service offerings to full service on-site deployments.
- Apstra Premium (future order)

THE RESULTS

- Insights provided by specialist team on current and future designs **established instant credibility**
- Extensive competitive evaluation proved **Juniper's superior ability to deliver** easy-to-deploy/operate, scalable networking driving faster time-to-value for AI clusters
- JVDs enable **quick spin up of AI-as-a-Service** to attract new clients

"We're able to deliver accessible, first-class AI transformation for our customers"

A perspective view of a modern server room. The room features rows of blue server racks filled with various electronic components. The floor is dark and reflective, with glowing green light strips running along the base of the racks. The ceiling has recessed lighting. The overall atmosphere is high-tech and futuristic.

Next Gen of AIDC

從早期採用者..到大眾市場

專有人工智慧

不斷演進的人工智慧解決方案

單一來源的 A100 和 H100 GPU

競爭激烈的 GPU 供應商市場不斷擴大

封閉式 Nvidia AI/ML PyTorch 框架

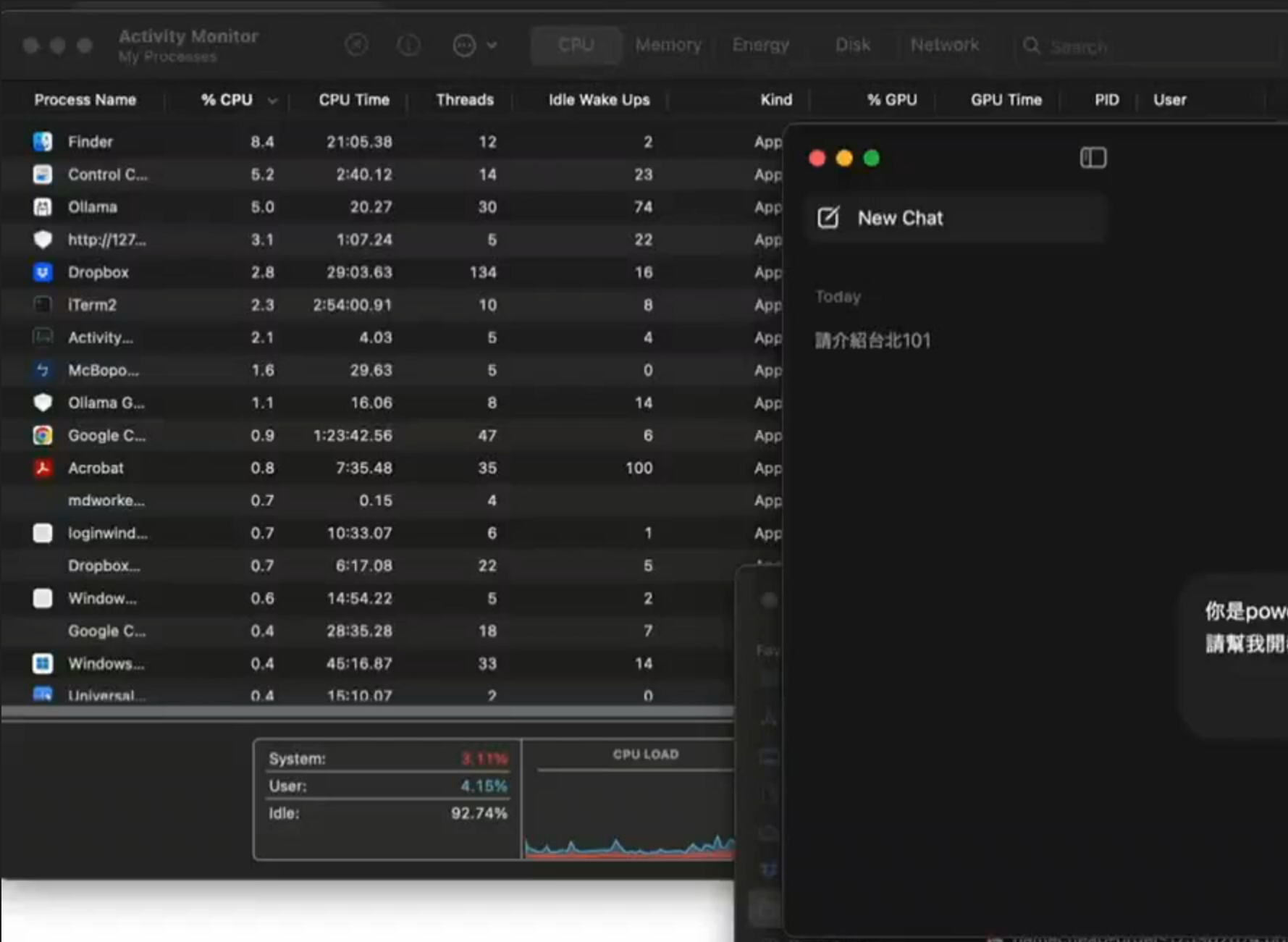
PyTorch 2.0 擴展了 GPU 支援和生態系統

緊耦合的專有 InfiniBand AI 網路

開放式乙太網 Fabric 可實現 Tb 速度 - “以太網絕對適用於人工智慧訓練”
SaaS 供應商網路運行與工程主管

單一供應商創新

行業驅動的創新, UEC



New Chat

Today

請介紹台北101



你是powershell專家

請幫我開發一個powershell,列出密碼在14天內要到期的使用者

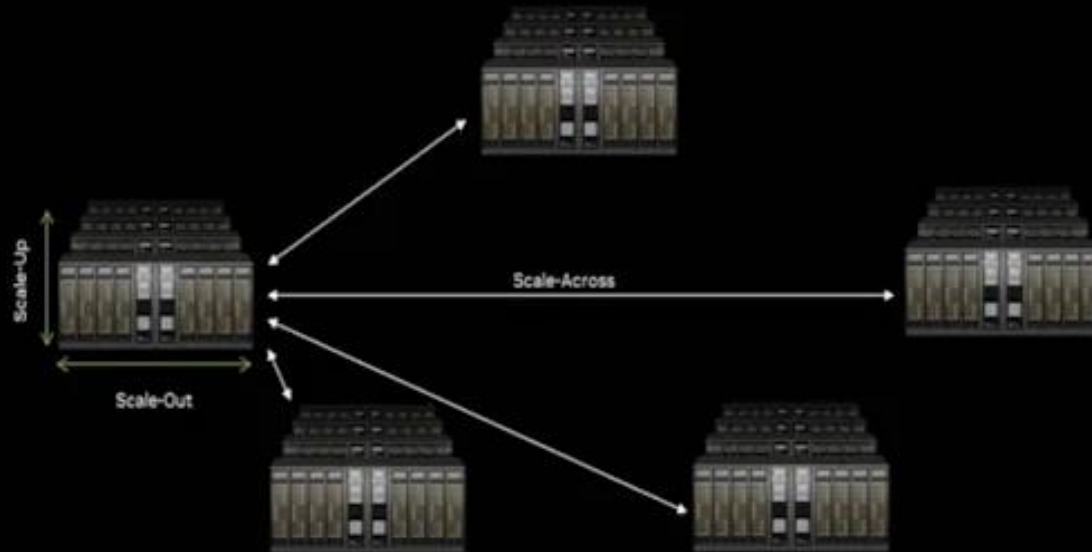
gpt-oss:20b



NVIDIA announce Scale-Across solution

NVIDIA Introduces Spectrum-XGS Ethernet to Connect Distributed Data Centers into Giga-Scale AI Super-Factories

Distributed AI between remote locations, overcoming power and physical limitations



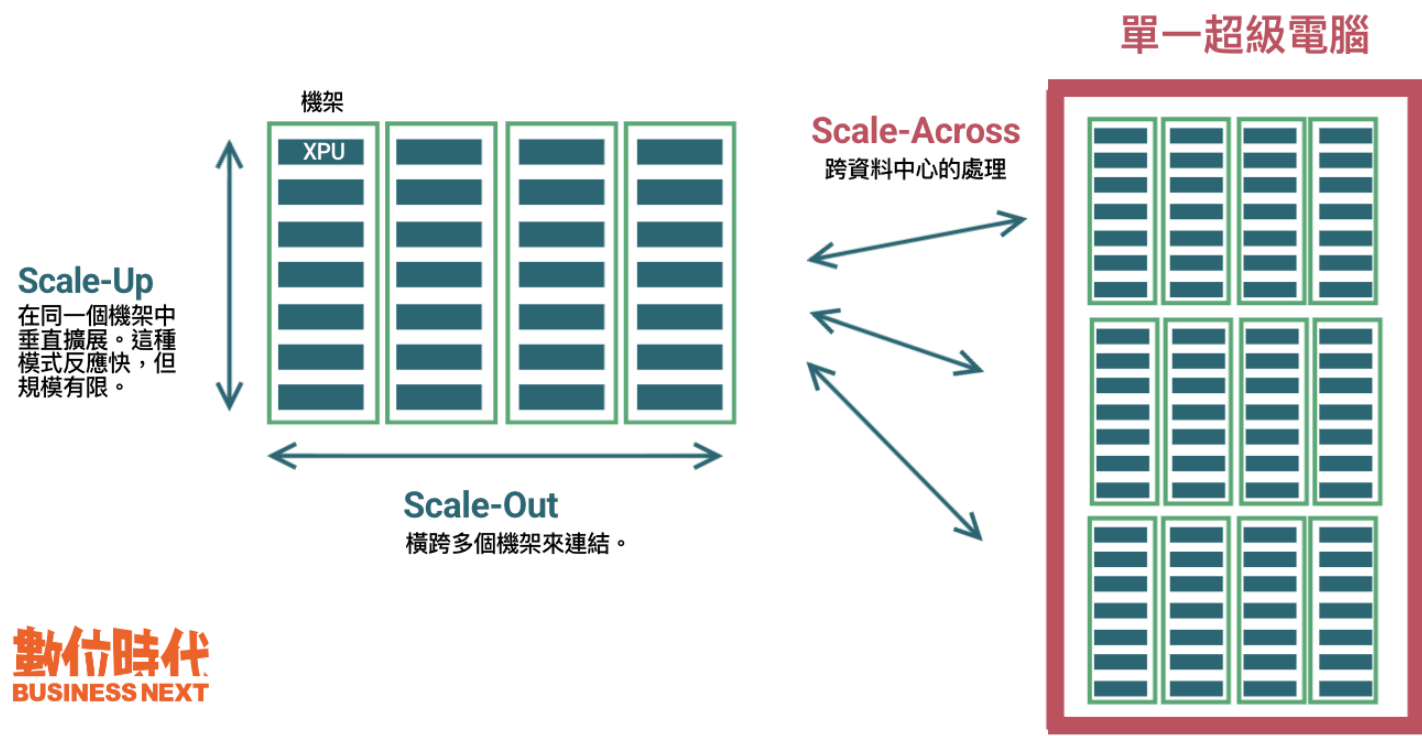
- Scaling AI beyond the data center requires new infrastructure
- Spectrum-XGS unifies multiple data centers into the world's largest supercomputers
- Integrating scale-out and scale-across
- Auto-adjust load balancing based on scale-across distance
- 1.9X higher NCCL multi-site performance

What is Scale Across

Scale-Across是什麼？

目前企業建構AI基礎設施時，主要採取垂直擴展（Scale-up）與水平擴展（Scale-out）兩種連結模式，輝達的「Scale-Across」，可讓多個資料中心透過新一代交換技術互聯，形成如同單一超級電腦的運算體系。

- Scale Across」與傳統的「Scale-Up」（單一機櫃/伺服器擴容）和「Scale-Out」（橫向擴展更多機櫃/增加伺服器）不同，強調的是跨資料中心的協同運算能力。
- 突破了單一資料中心的規模及電力限制，能把多個城市、甚至不同國家的運算資源整合起來共同執行AI和大型運算任務。



數位時代
BUSINESS NEXT

Share

made with infogram



 **nVIDIA.**

DGX SPARK



Technical Specifications*

Architecture NVIDIA Grace Blackwell

GPU NVIDIA Blackwell Architecture

CPU 20 core Arm, 10 Cortex-X925

+ 10 Cortex-A725 Arm

CUDA Cores NVIDIA Blackwell Generation

Tensor Cores 5th Generation

RT Cores 4th Generation

Tensor Performance 1 PFLOP

System Memory 128 GB LPDDR5x, unified
system memory

Memory Interface 256-bit

Memory Bandwidth 273 GB/s

Storage 1 or 4 TB NVME.M2 with self-encryption

USB 4x USB TypeC

Ethernet 1x RJ-45 connector 10 GbE

NIC ConnectX-7 Smart NIC

Wi-Fi WiFi 7

Bluetooth BT 5.3 w/LE

Audio-output HDMI multichannel audio output

Power Consumption TBD

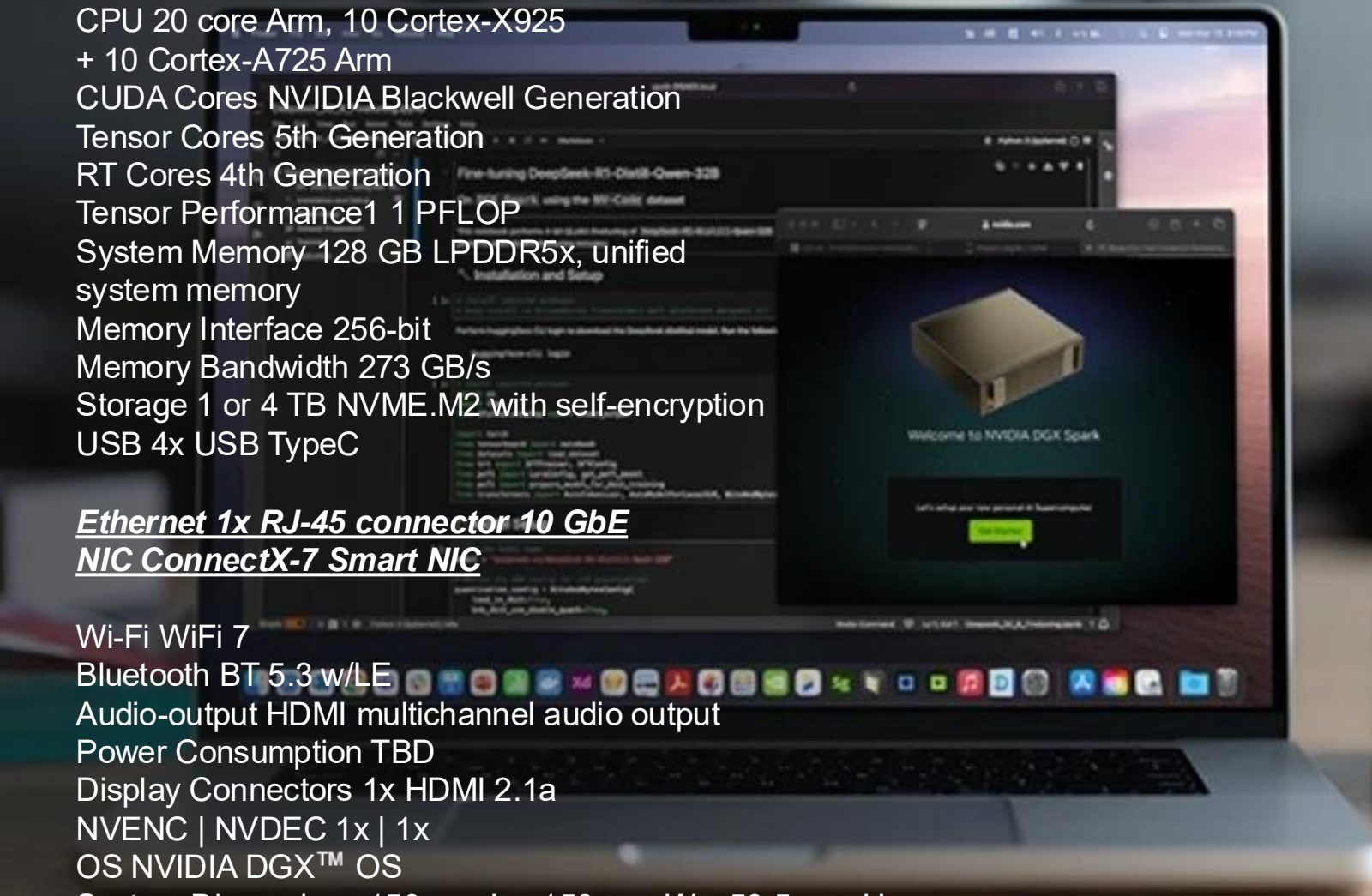
Display Connectors 1x HDMI 2.1a

NVENC | NVDEC 1x | 1x

OS NVIDIA DGX™ OS

System Dimensions 150 mm L x 150 mm W x 50.5 mm H

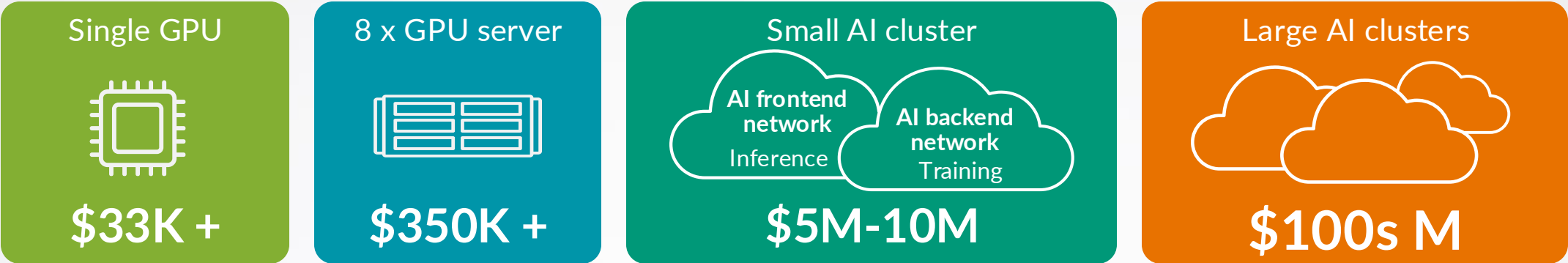
System Weight 1.2 kg





Shaping the Next Step Forward

GPUs are expensive and scarce



Data center CapEx comparison

	Traditional DC	AI Training DC
Compute	55%	80%
Storage	35%	14%
Network	10%	6%

Backend AI DC switching TAM:
\$3B (2023)
65% CAGR (2022-2027)

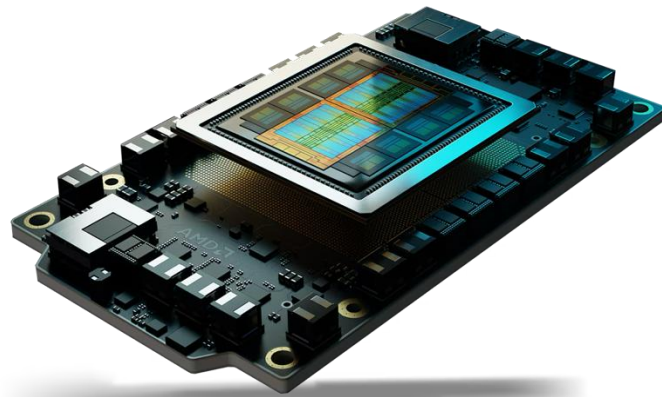
...the network is critical

Meta claimed last year that 33% of elapsed time in AI/ML is spent waiting for the network

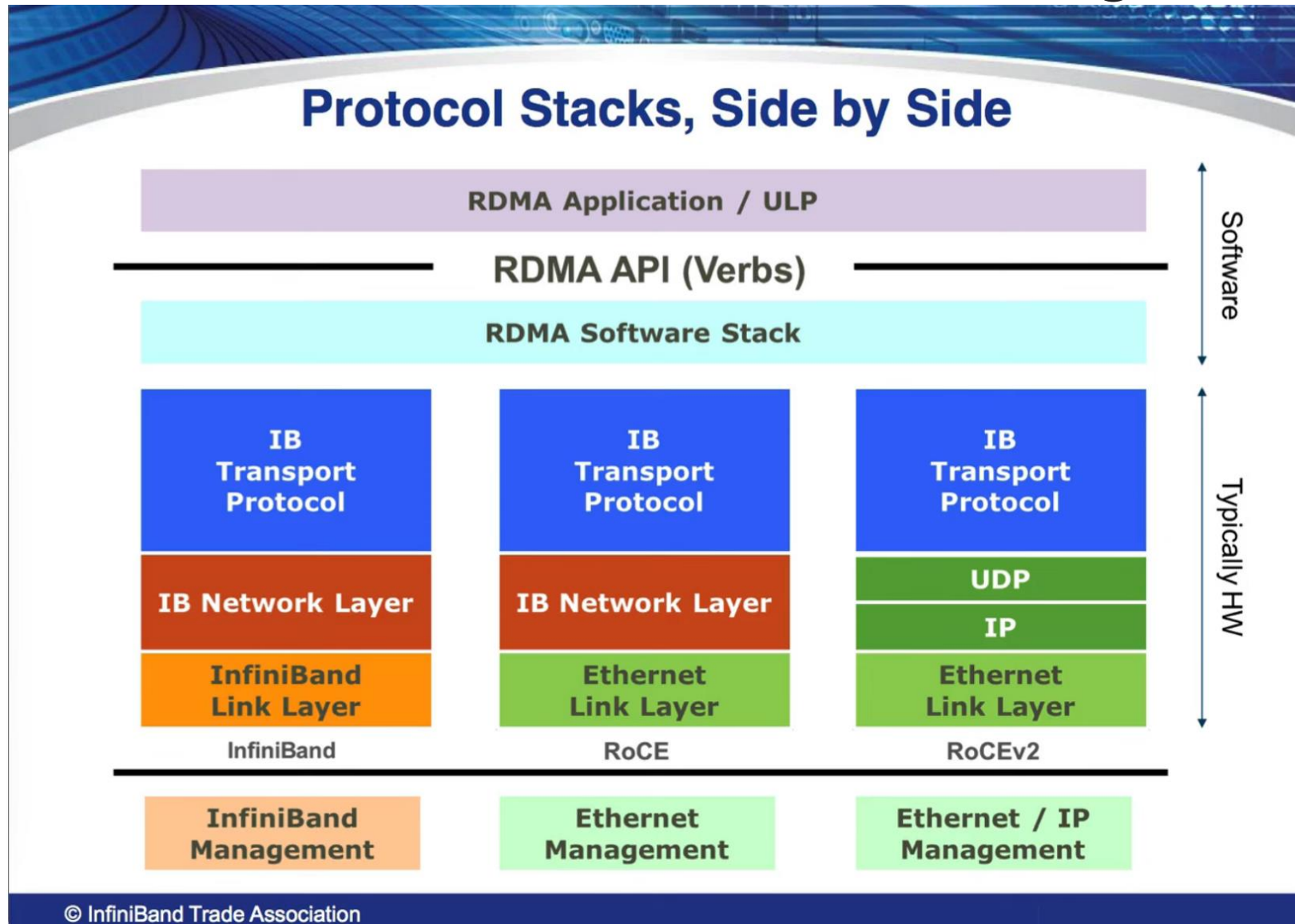
Source: Dell'Oro

Step Forward

- Market Trend of GPU/CPU/Memory resource
 - NVIDIA is priced beyond the reach of many users.
 - Most models aren't tied to AMD or Intel GPUs; in fact, multi-vendor deployments are already common.
 - The majority of AIDCs now use multi-vendor computing architectures



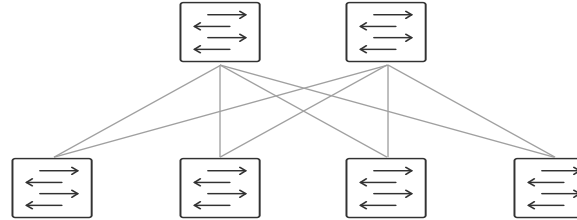
What is RoCEv2 (RDMA over Converged Ethernet v2)



<https://www.youtube.com/watch?v=8kTAXhujn08&t=143s>

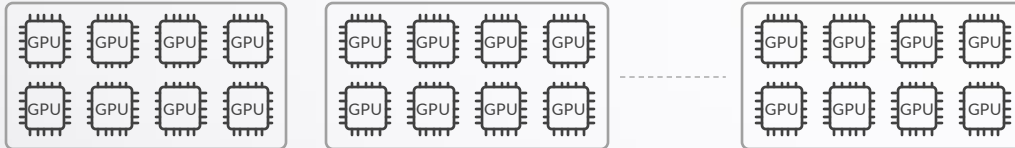
AI Data Center fabrics

Front-end Fabric / Inference Fabric

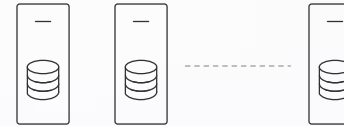


- User to workload connectivity
- Typically, 10/25/100G
- Ethernet

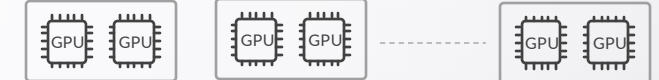
Training nodes



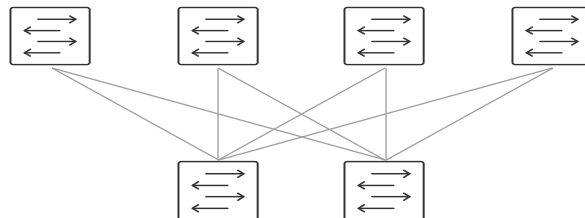
Storage nodes



Inference nodes

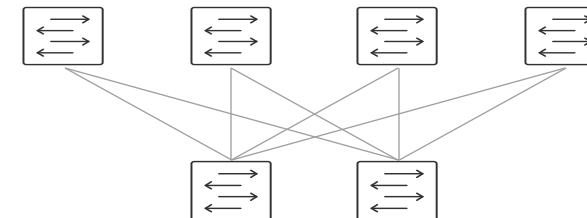


Training Fabric (Backend)



- Inter-node GPU communication
- Typically, 400G
- InfiniBand or Ethernet

Storage Fabric (Backend)



- GPU-to-Storage communication (Training & Inference)
- Typically, 100/200G
- InfiniBand or Ethernet

Ultra Ethernet Consortium



Formed to create a new communication stack for Ethernet that is high-performance and open:

Our Mission

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

Adopting a clean slate approach to developing a complete communications stack for AI and HPC. A 1.0 Spec is planned for 2025 publication. UET is more than just a RoCEv2 replacement.

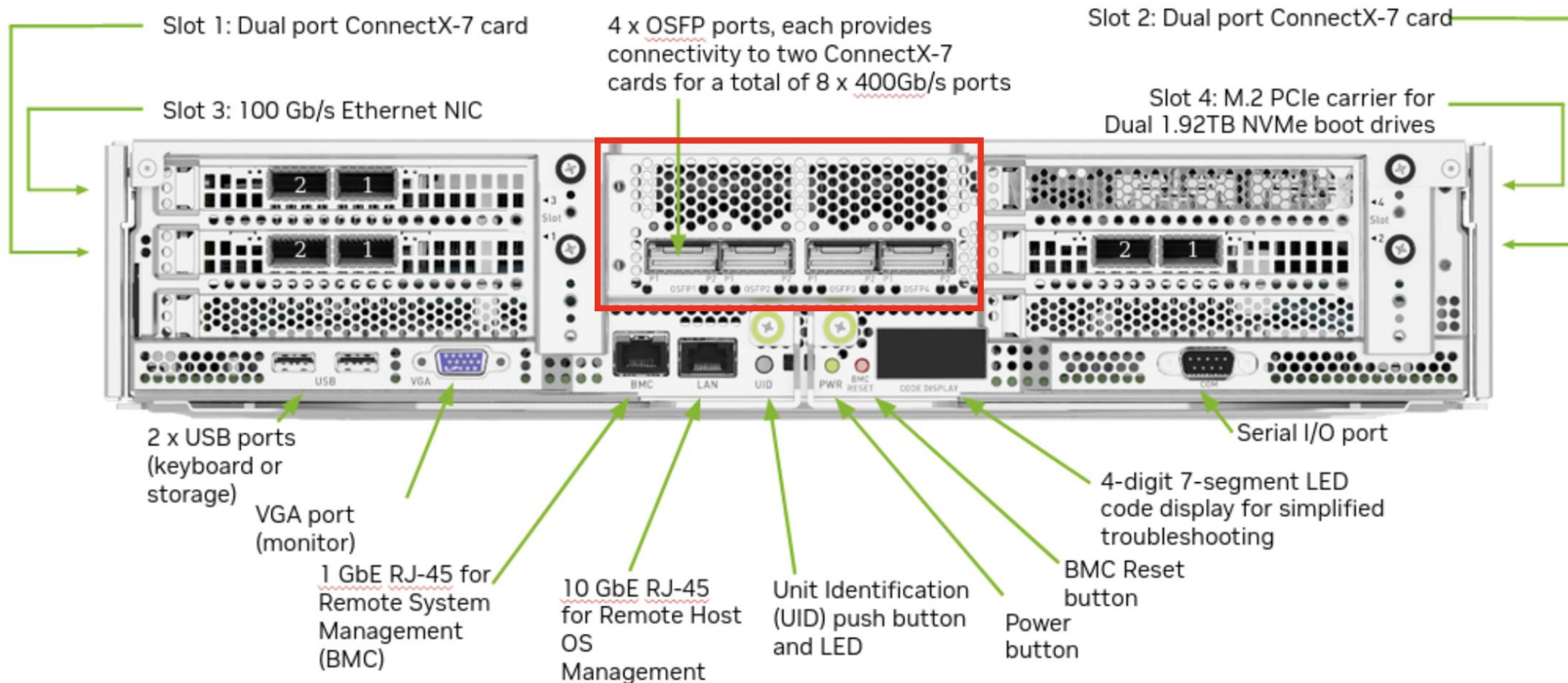
Juniper's view:

- Juniper is a General Member of UEC since Nov 2023
- A full spec for UET will take time to develop and even more time for adoption by hardware vendors
- A full communication stack will take significant time to develop so while we intend to participate in UEC, we will also be focused on tuning and optimizations for RoCEv2 which is heavily used in AI/ML today

<https://www.juniper.net/us/en/the-feed/topics/ai-and-machine-learning/juniper-networks-ai-data-center-and-ultra-ethernet-consortium-uec.html>

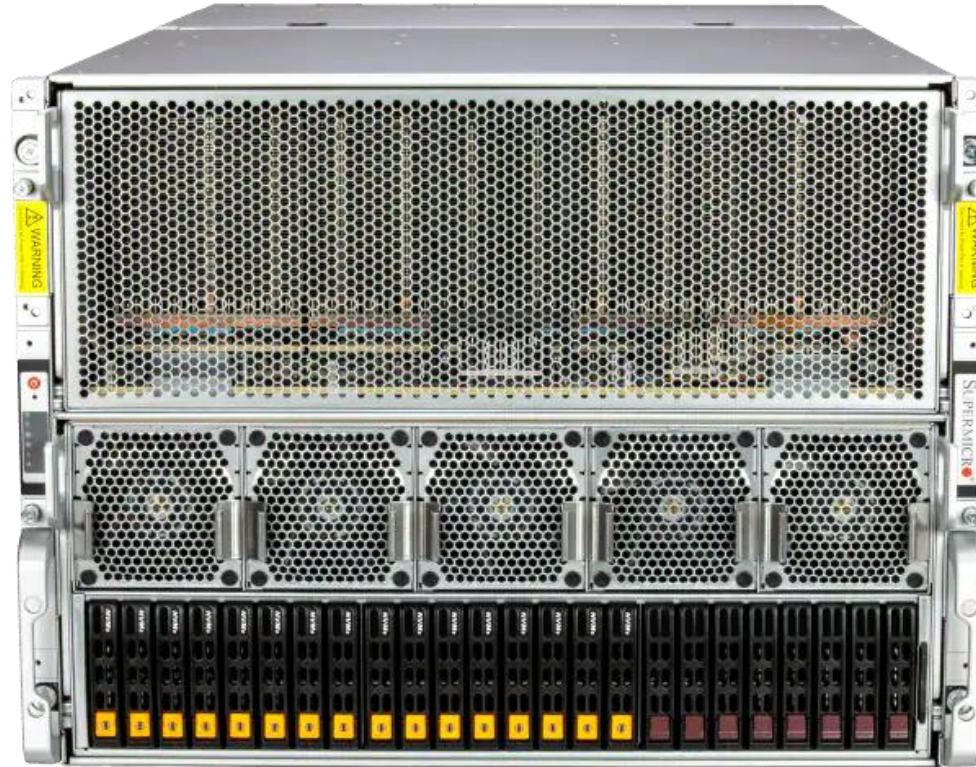
DGX H100 – Network Ports

DGX H100 - NVIDIA Server with GPUs to accelerate deep learning applications



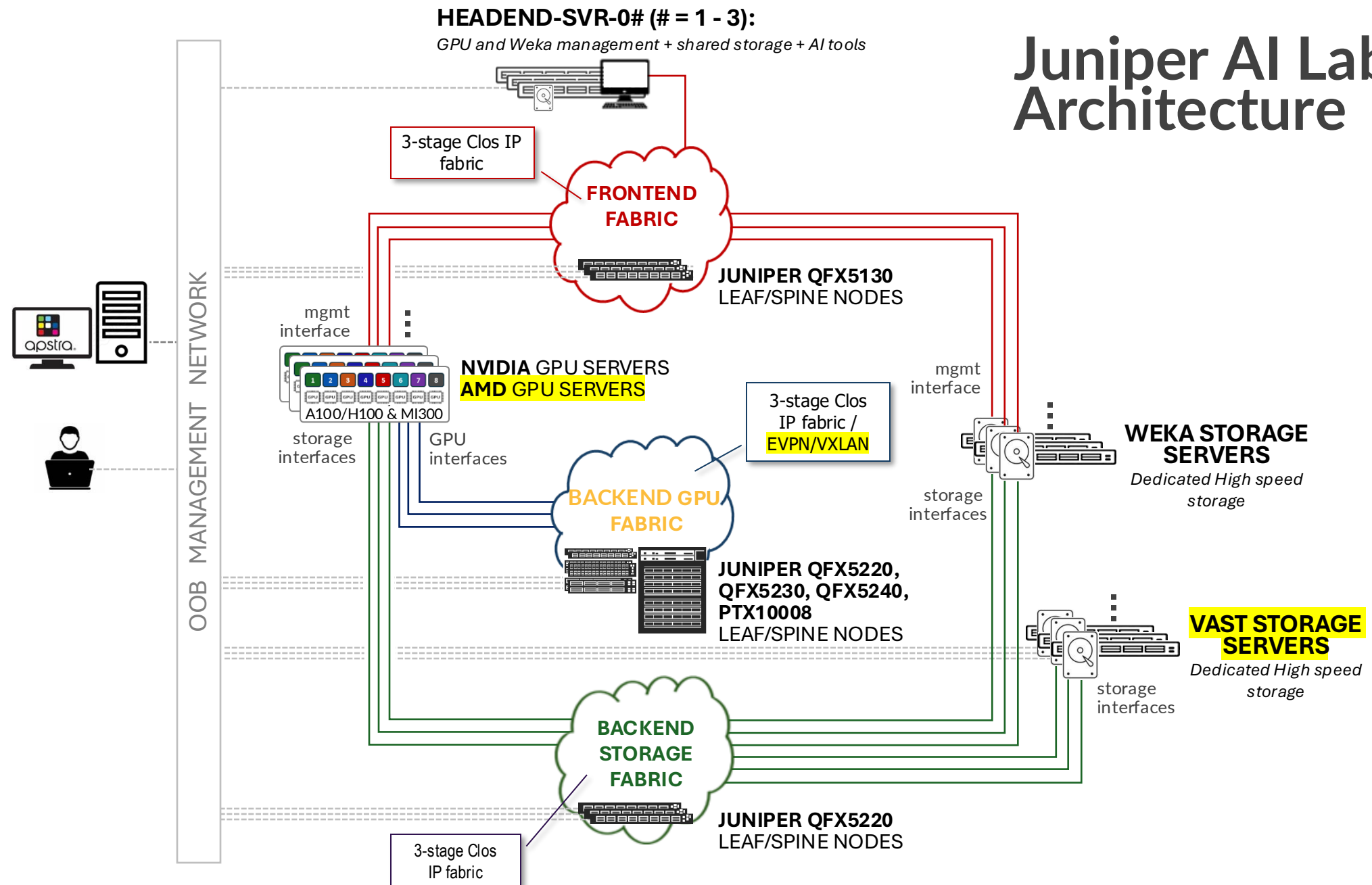
AMD Instinct MI300x Systems

- 8 Rack Unit
- 4 x AMD EPYC/Genoa 9554 CPU
- 8 x AMD Instinct MI300x GPU / 192GB
- 2TB Memory (20x96GB)
- 4 x Samsung 1.92TB SATA
- 16 x Samsung 3.8TB NVMe
- 8 x Connect-X 7 400G QSFP112 or Broadcom BCM57608 400G (thor 2)
- 4 x Connect-X 7 200G dual port QSFP112
- Power consumption: 9,600W / system
- BTU: 32,675 / system



- 2 x Supermicro AS-8125GS-TNMR2 Dual AMD EPYC 8U GPU SuperServer
- 2 x Dell PowerEdge XE9680 6U GPU Rack Server

Juniper AI Lab Architecture



Open Ethernet offer best TCO with performance



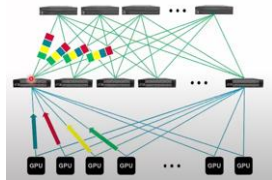
InfiniBand

Proprietary Mellanox NICs and Switches



AI Optimized Ethernet

*Shallow Buffer, Deep Buffer
across leaf and spine*



Scheduled Fabric

Cell or Ethernet-based disaggregated chassis

Vendor Lock-In

Yes

No

Yes

Operational
Consistency

No

Yes

No

Cost

Higher

Lower

Higher

Scale

Central/Limited

Distributed/Higher

Central/Limited

Performance
for AI workloads

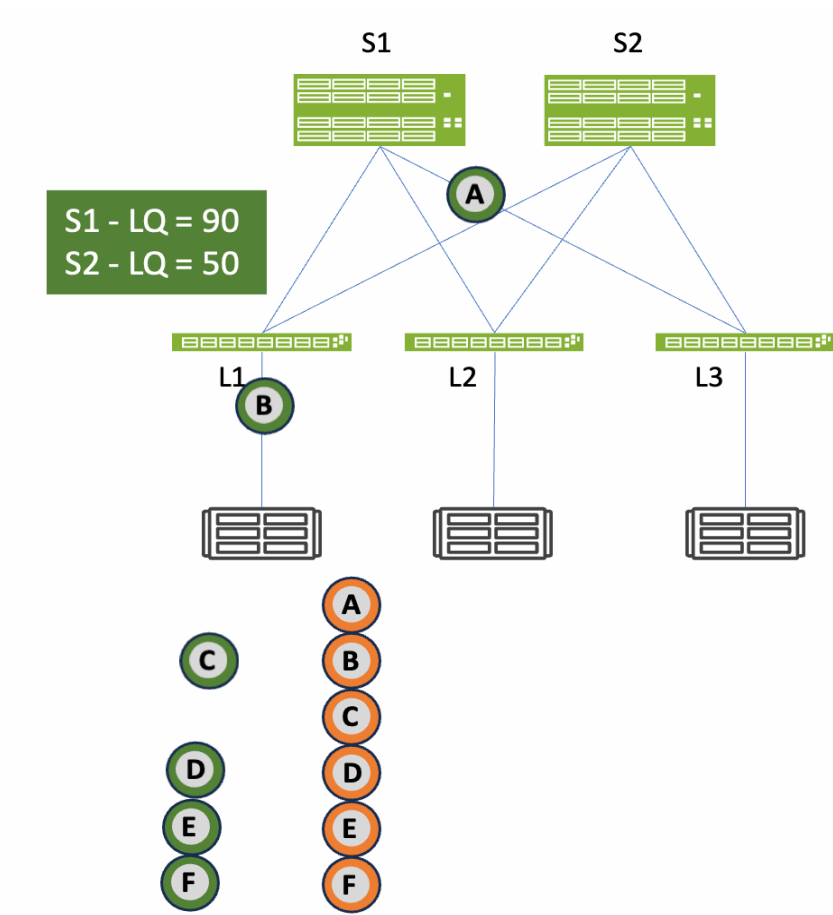
Yes

Yes

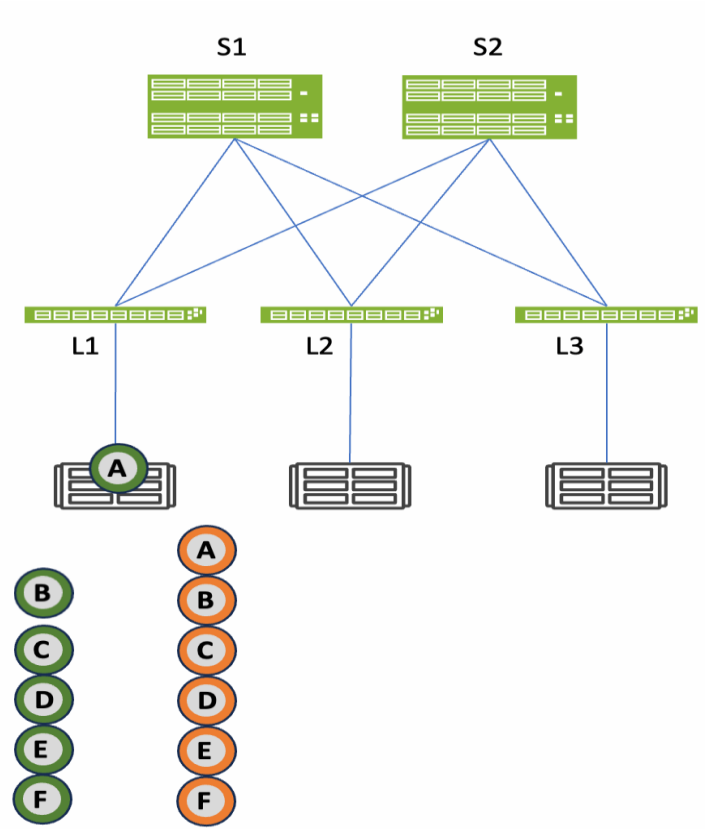
Yes

Network-based Dynamic Load Balancing (DLB)

DLB flowlet mode



DLB packet mode



Apstra 服務可視化的效益

Benefits

- planning/ bandwidth
- verify firewall rule
- Geo IP
- DDoS Attack
- troubleshooting
- Service awareness



支援NG-AI的網路設備應為考量重點

維運

ANSIBLE

Terraform



融合 AI NetOps

一致的人工智慧平臺 NetOps 工作流程和自動化，提供操作簡單性、速度和可靠性

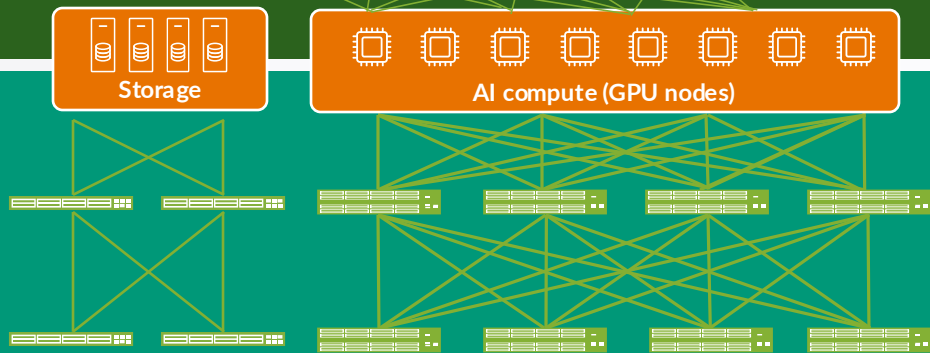
前端



100G/400G/800G 乙太網矩陣

QFX5120, QFX5130 交換機為開放式乙太網 Fabric 提供最佳性價比

後端



RDMA over Converged Ethernet

GPU 高效人工智慧基礎設施

新型高密度 400G/800G PTX, QFX 交換機為開放式乙太網 Fabric 提供最高容量和規模

IBN + AIOps 配有人工智慧擴展和先進的流量管理，可提供靈活性和更高的經濟性



Thank you

JUNIPER[®]
NETWORKS

Driven by
Experience[™]