

Building AI Network Fabrics

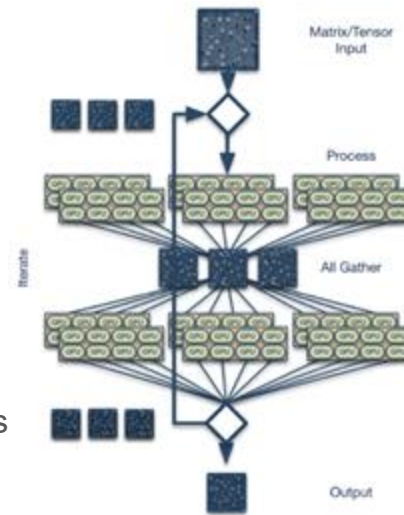
Agenda

1. Why is networking for AI different than traditional networking?
2. What is the Arista strategy for AI networking ?
3. What do these networks look like ?
4. What are the Arista platforms for AI networking ?

Why is networking for AI different than traditional networking?

AI workloads main characteristics

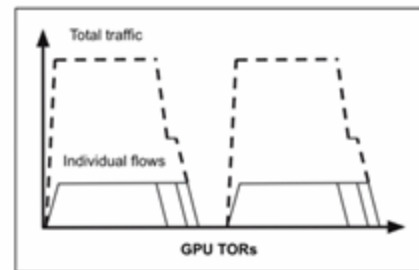
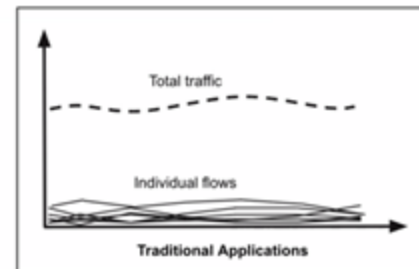
- It requires specialized Hardware and Software
 - XPU (GPU, ICU, TPU ...)
 - Library and development frameworks
- Principles
 - AI applications are “optimization” applications (with potentially billions of parameters to look after ...). These are executed on XPUs.
 - These XPU’s execute algorithms/calculations in parallel and then share results amongst each other (“sharding” vs “replicating”).
 - » There are several ways of doing this, the commonly used term is “Collective Operations”
 - All XPU’s have to wait until all XPU’s have shared their messages and passed on results, the workload is stalled until this happens, commonly used term is “Barrier”
- Job Completion Time (JCT) is the key metric



What are the consequences ?

AI traffic characteristics

- This is not TCP but RDMA (Remote Direct Memory Access)
- RDMA NICs push at line rate of 100G/200G/400G
 - Traffic has a small number of large flows
 - >> Very poor entropy
- The effects from the collective communications
 - Bursty traffic
 - Latency accumulation for “ring-based” collectives
 - Congestion in “tree-based” collectives
 - Fails “together”
- Current DC’s are on demand, generalist applications and can recover easily from failure, in the AI world failures affect performance
 - Typical Data Center applications can easily recover from failure, in AI any type of failure will affect performance. Then uses checkpoints is critical in AI applications.



Arista Strategy for AI networking

Arista Next Generation AI Etherlink Portfolio

3

New AI Optimized 800GbE Systems
Maximizing choice



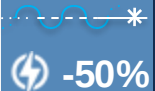
Accelerator, Workload and NIC
Agnostic



Best in Class Performance
for AI Workloads



Latest Generation 5nm Geometry
2x Density and 25% Less Power



>50% Power Reduction with
LPO and Extended DAC Cables



Scaling AI from 10s to > 100,000 Accelerators

Key architecture building blocks

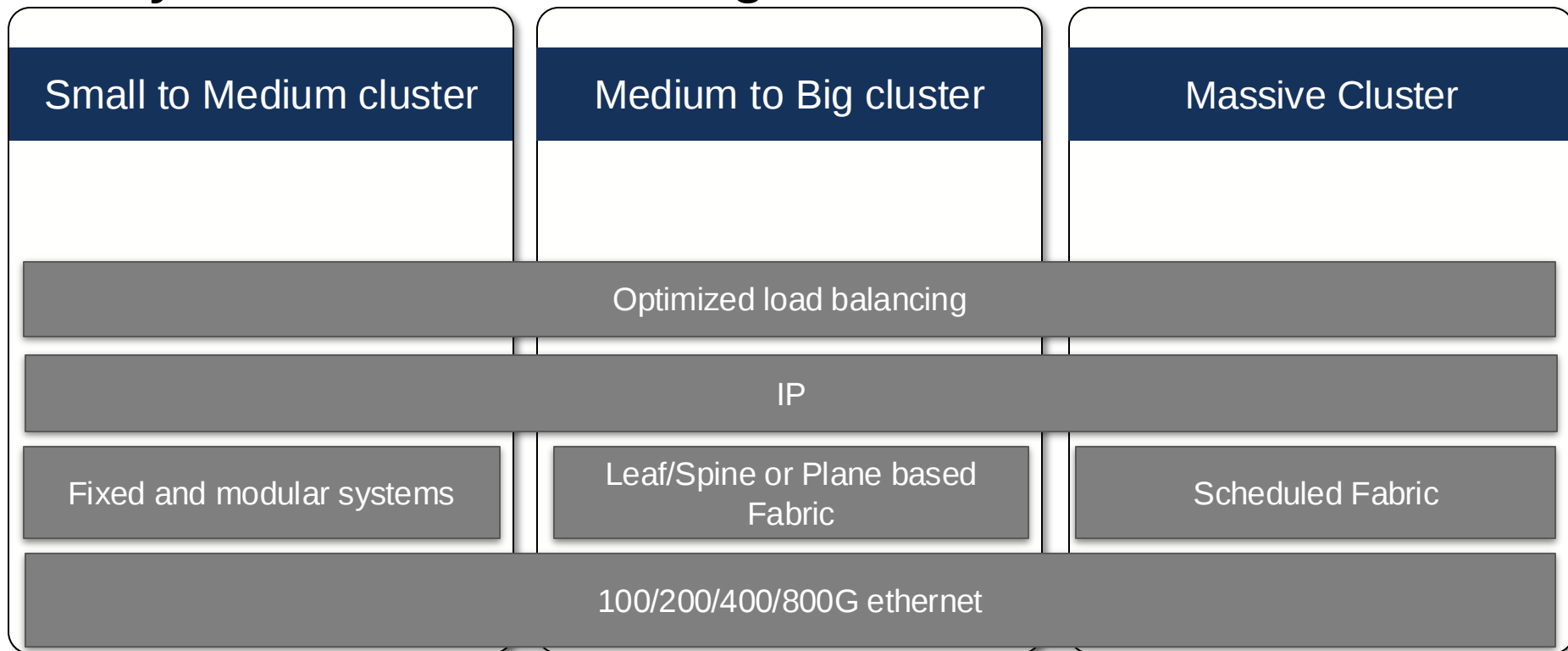
Small to Medium cluster

Medium to Big cluster

Massive Cluster

Accelerator & NIC Agnostic, Open Standards, Smart AI Features

Key architecture building blocks



Accelerator & NIC Agnostic, Open Standards, Smart AI Features

Arista AI Portfolio

Small to Medium cluster



Fixed & Modular Systems

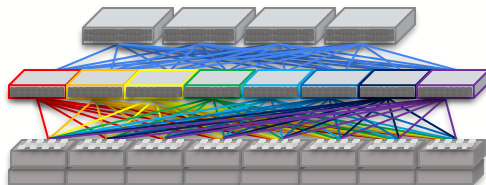
51.2T ➡ 460.8T

64-576 x 800G Nodes

128-1152 x 400G Nodes

Lowest Cost, Power &
Complexity

Medium to Big cluster



2- or 3-Tier Leaf-Spine or Plane-based

Multi-Petabit Bisectonal B/W

DLB, PFC, ECN,

Tens of Thousands of Nodes

Massive Cluster



Single Hop Interconnect

Fully Scheduled Lossless

100% Efficient, Cell Spraying

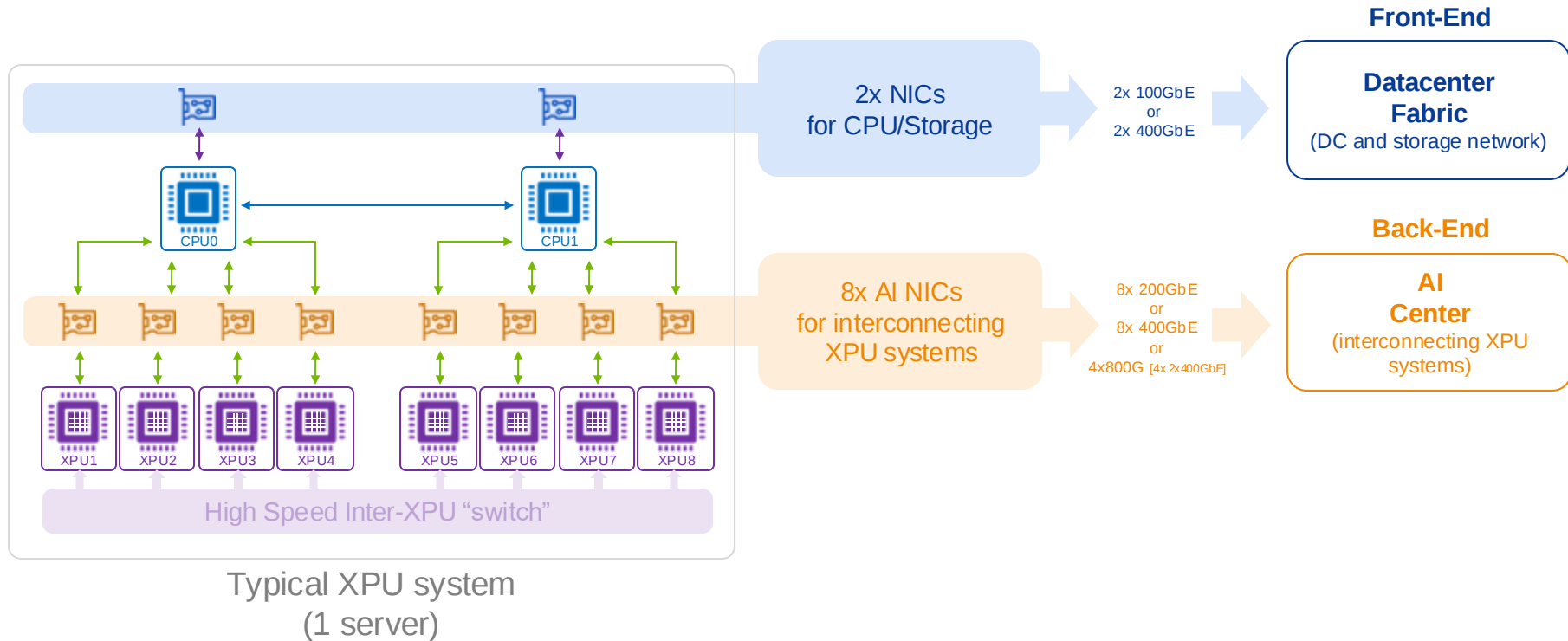
Multi-Petabit Bisectonal B/W

Tens of Thousands of Nodes

Accelerator & NIC Agnostic, Open Standards, Smart AI Features

What do these AI networks look like ?

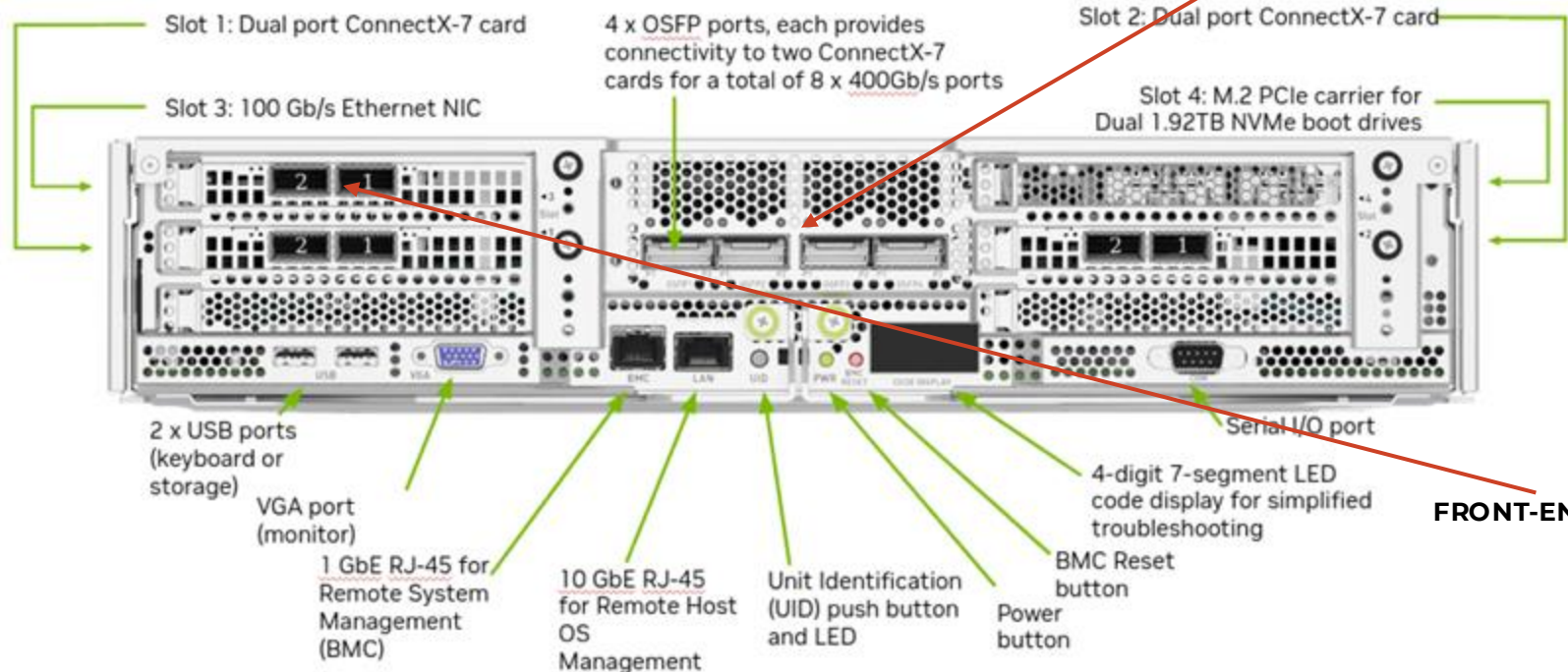
Where to connect XPU systems ?



NVIDIA NCIS - H100 motherboard

Here is an image that shows the motherboard connections and controls in a DGX H100 system.

BACK-END ports



FRONT-END ports

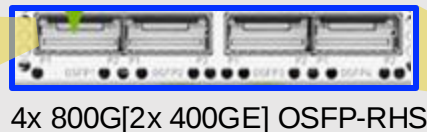
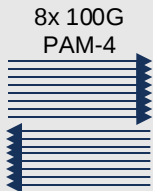
What are the Arista platforms for AI networking ?

Connecting High Speed NIC / XPU Port Types

NIC/XPU Port Form Factor

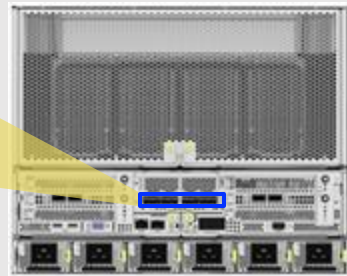
Example HW (not exhaustive, examples only)

800G OSFP-RHS
(2x 400GE)
“flat top twin port OSFP”

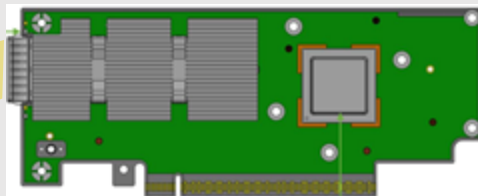
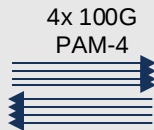


4x 800G[2x 400GE] OSFP-RHS

NVIDIA DGX H100

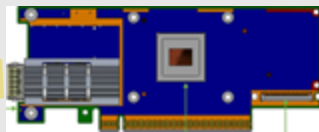
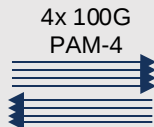


400G OSFP-RHS
“flat top single port OSFP”



NVIDIA ConnectX-7

400G QSFP112



NVIDIA ConnectX-7



NVIDIA Bluefield 3



Broadcom Thor2

Platforms for AI Networking - 2024

Fixed

Low Latency, Fixed Systems



7060X6-32E – 32x 800G



7060X6-64E – 64 x 800G



7060X Series
25.6T - 51.2T Systems

Modular

High Density, VOQ, Cell Based



36 x 800G Linecard



7800R Series 4 to 16 Slots
Up to 460T modular System

Next Generation – 7060X6 Series for 800G



64 x 800G AI Optimized Leaf



32 x 800G AI Optimized Leaf

High Performance and Radix:

- Up to 64 x 800G wire-speed ports
- 51.2T and 25.6T Bandwidth

Advanced AI Features:

- Advanced load balancing - DLB, RDMA Aware LB
- SSU, Network Telemetry, DCQCN
- EVPN VXLAN for multi-tenancy
- LPO Support

2-tier scaling to thousands of ports

Next Generation 7800R4 Series



NEW: 36 x 800G AI Optimized Line Card

High Performance Modular System:

- Up to 576 x 800G / 1152 x 400G
- Non-blocking 460 Tbps (Full Duplex)
- Foundation for Scale Out AI:
 - AI optimized pipeline
 - Lossless fully scheduled VOQ
 - 100% fair cell-based fabric
 - Non-blocking architecture
 - Integrated over-provisioning
 - Deep buffering for sustained traffic
- Ideal for single tier clusters or 2-tier scaling to tens of thousands of ports

AI Center

AI Center

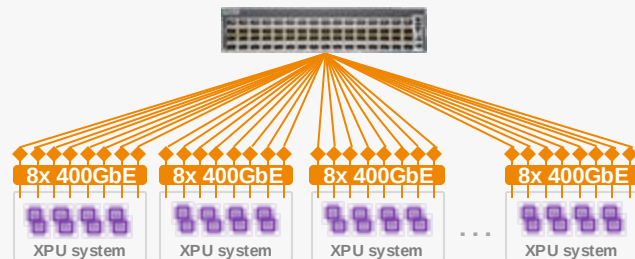
Key requirements

- 400GbE access ports
- No oversubscription
- Optimized flow distribution
- Lossless
- Advanced Telemetry

Key Variables

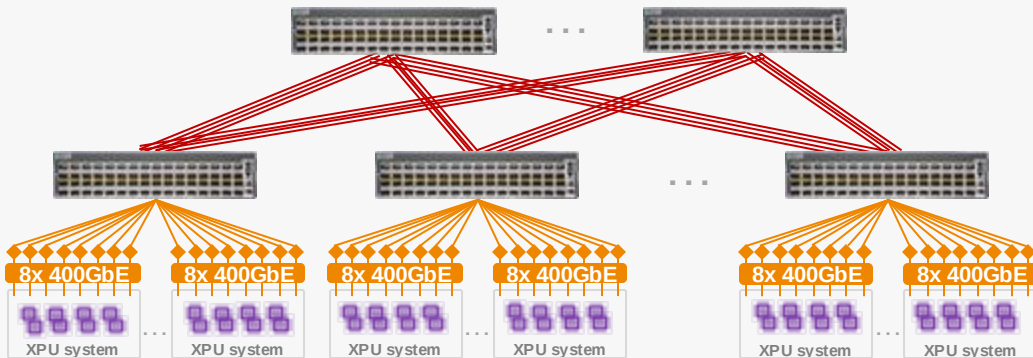
- Total # of AI NIC ports
- AI NIC Transceivers cap.
- Rack physical layout and fiber plant
- Cost

Single-tiered AI Fabric



Small and Moderate AI applications (10s and 1k of xPUs)

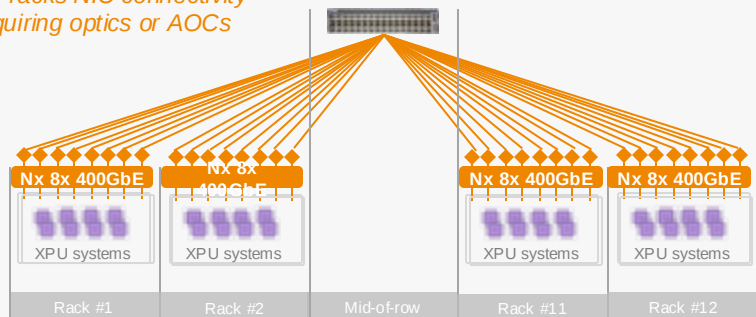
Multi-tiered AI Fabric



Large AI applications (100ks of xPUs)

Single-tiered fixed

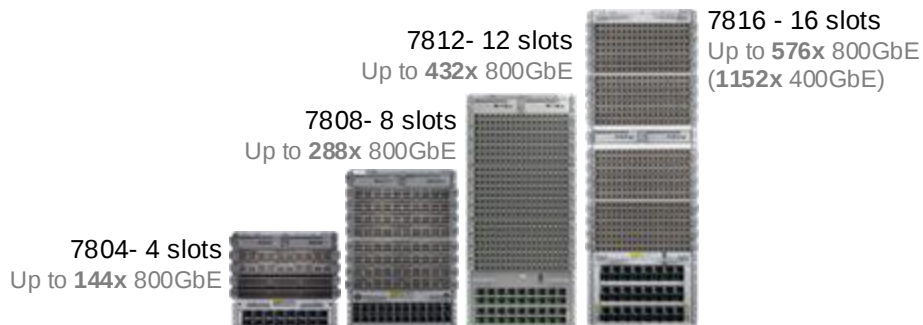
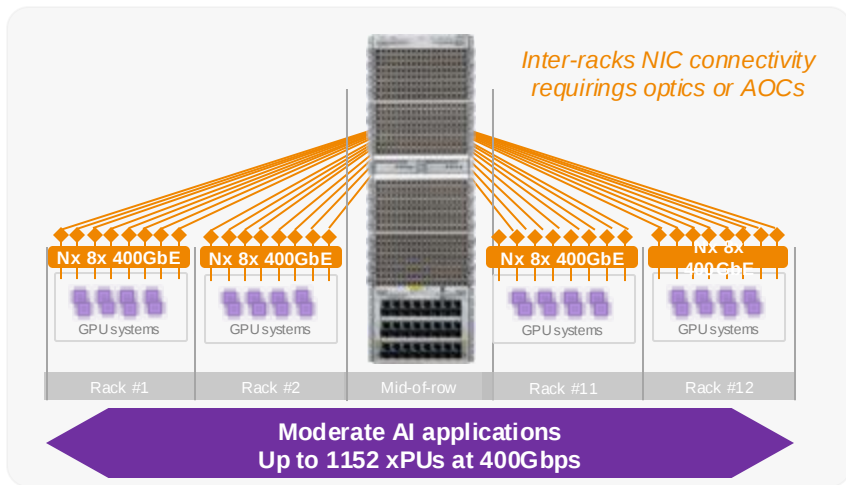
*Inter-racks NIC connectivity
requiring optics or AOCs*



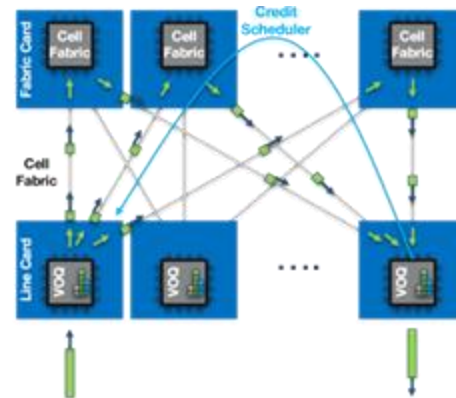
7060X6-64E
64 port 800G OSFP

- Fixed Configuration Switch
 - Up to 64x 800G
- No flow collisions
 - Single-asic line-rate forwarding
- ECN and/or PFC to handle incasts
 - Low buffers - Requires tuning

Single-tiered modular

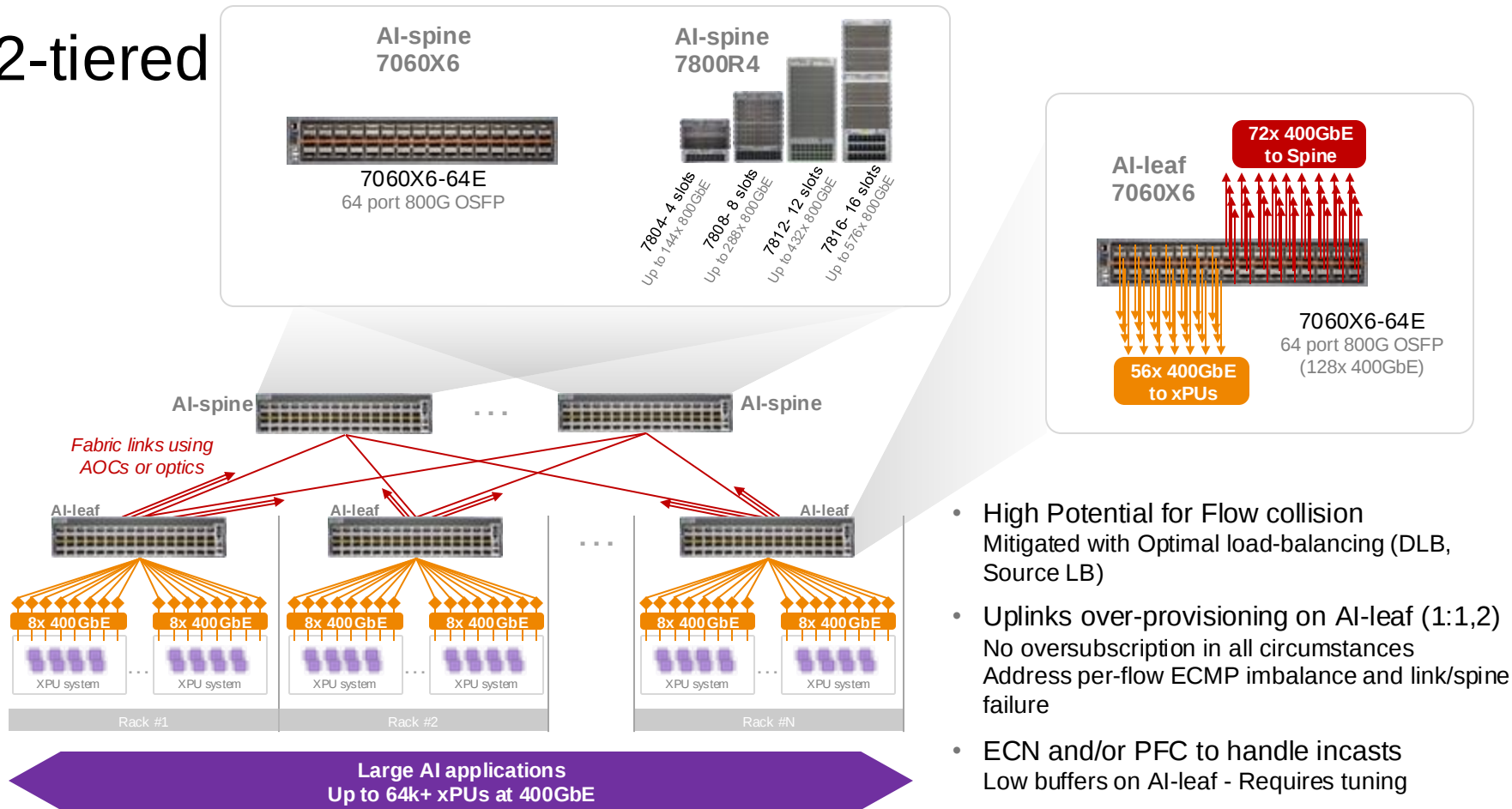


- 7800R4 Modular chassis offering high port density
4, 8, 12 or 16 slots / 36x 800GbE Linecards
- Non-blocking distributed forwarding
Leaf (Linecards) & Spine (Fabrics) in a single chassis
- No flow collisions between line card and fabric
Scheduled Cell-based Fabric
Built in overprovisioning between line card and fabric
100% Fair and Efficient Load Balancing within the chassis



- High Availability
Fabric, fan, power supply, sup redundancy
- ECN and/or PFC to handle incasts
Deep buffers - Requires minimal tuning

2-tiered



Scaling out the R4 AI Spine

7800R4 AI Spine



Lossless Fully Scheduled VOQ
100% Fair And Efficient Traffic Spraying
Integrated Redundancy and Resilience
Optimized Pipeline for AI Workloads

Scale Out

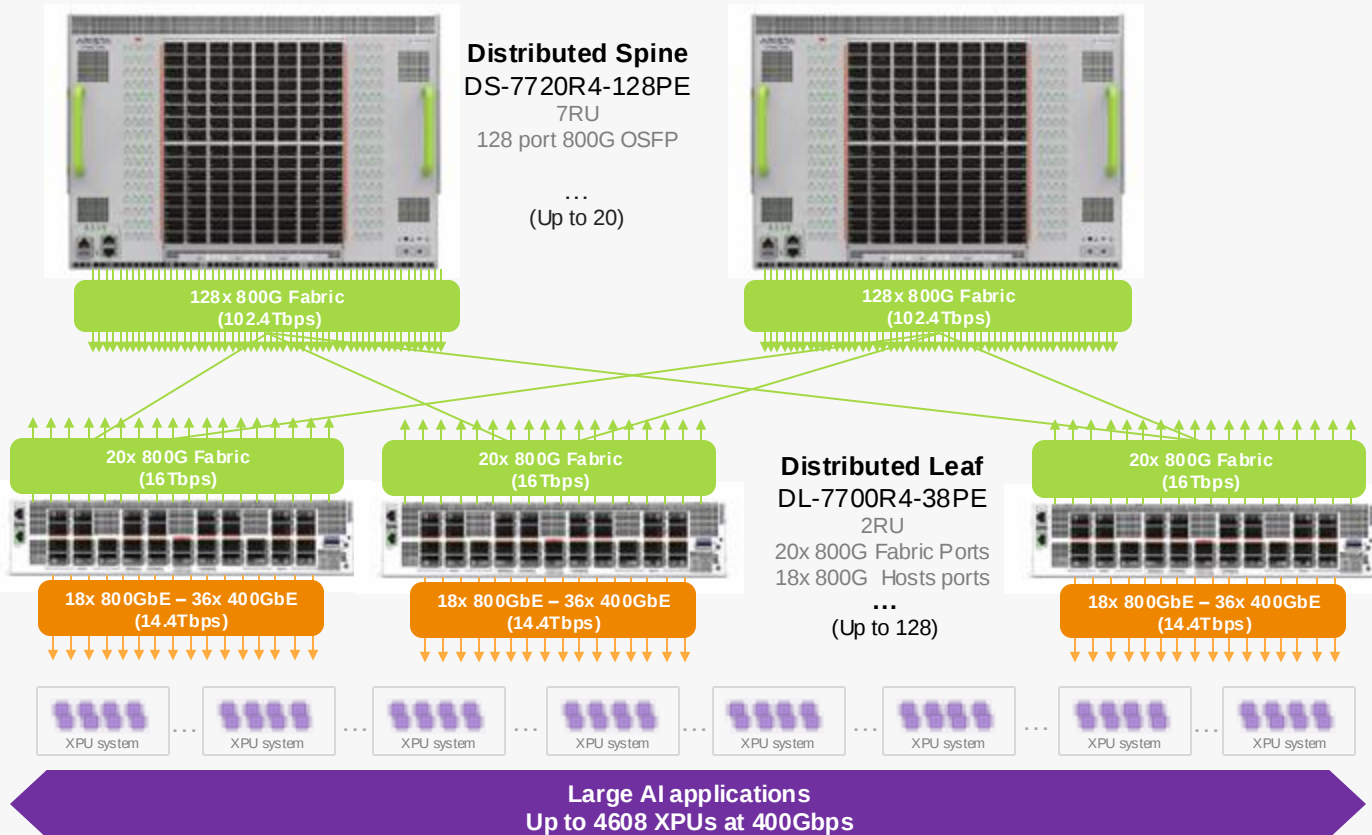
7700R4 Distributed Etherlink Switch



Lossless Fully Scheduled VOQ
100% Fair And Efficient Traffic Spraying
Integrated Redundancy and Resilience
Optimized Pipeline for AI Workloads

Common Architecture – Re-packaged for AI Scale Out

Distributed Etherlink Switch – 4k+ 400GbE XPU ports

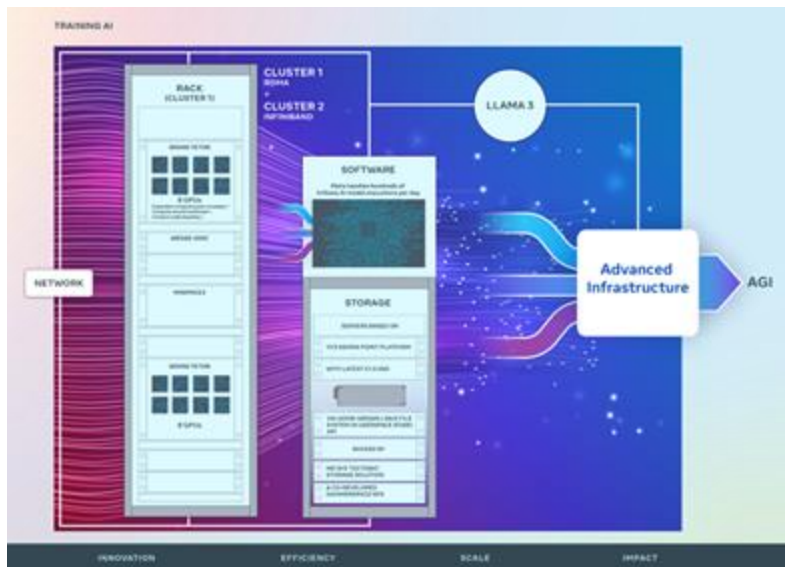


- Single-tier distributed switching system
 - Single logical switch
 - No tuning
- No flow collisions
 - 100% efficient VOQ
 - Cell Architecture
- Rich EOS and CloudVision
 - Single point of management
 - Independent device upgrade and replacement

Ethernet AI success story

Ethernet AI success at Meta

.. we have successfully used both **RoCE** and InfiniBand clusters for large, GenAI workloads (including our ongoing training of Llama 3 on our **RoCE** cluster) **without** any network bottlenecks..



- 2 Clusters with 24,576 NVIDIA Tensor Core H100 XPU each:
 - one based on NVIDIA IB
 - one based on Ethernet (Arista 7800 + Wedge and Minipack aka Arista 7388X5)
- Similar Results, Similar performance, Similar End to End Support
- Only one technology is open and not only supported by one vendor. **Ethernet!**

<https://engineering.fb.com/2024/03/12/data-center-engineering/building-metas-genai-infrastructure/>

Conclusion

Why Arista is leading in AI networking?

Software Quality

Lowest Regression incidents and CVEs in the Industry

Software Extensibility

Customizable, Open and Interoperable

Software Manageability

CloudVision, SSU, BGP Maintenance Mode

AI
Center

Topology options

Layer 3 Leaf/Spine and high density Leaf only options

Hardware Choices

Most extensive options of form factors in the industry

Hardware Reliability

Highest MTBF with the lowest power consumption

XPU agnostic, best performance, vetted at Scale

Already deployed in XPU clusters of 8K, 16K, 32K and growing

Resources

[Arista AI White Paper](#)

[Arista AI Networking](#)

[Arista in RDMA Networks](#)

Tech Library

[RoCEv2 Deployment Guide](#)



[Arista 7800R3 Datasheet](#)

[Arista 7388X5 Datasheet](#)

[Arista 7060DX5 Datasheet](#)

[Arista 7060X6 Datasheet](#)

[Arista 800GbE FAQ](#)

[Arista 400GbE
FAQ](#)

[Arista 200GbE FAQ](#)



ARISTA

Thank You

www.arista.com