

AI FABRIC

The GPU servers arrived; the switching didn't keep up. Training jobs stall on communication, storage can't feed the cluster, and the next expansion doubles the problem. You need a lossless Ethernet fabric engineered for the cluster you actually run.

PRODUCT FAMILIES

Leaf / spine

AS7726-32X (100G) · DCS-series 400G platforms — 400/800G SKUs confirmed per cluster size at design time

NOS

SONiC (Enterprise or Community)

Optics / DAC

ET7402-SR4/CWDM4 (100G) · ET7302 (25G server) — 800G optics specified at design from Edgecore's transceiver line

SIZING & TIMELINE

S ≤64 GPUs, 1–4 racks — 3–5 weeks

M ≤512 GPUs — 6–10 weeks

L 512+ GPUs, multi-pod — 10–16 weeks, phased

SCOPE OF DELIVERY

- 01 Cluster workload profile and fabric sizing
- 02 Engineered design + BOM
- 03 Hardware supply
- 04 Fabric bring-up incl. RoCE tuning
- 05 Acceptance: measured all-to-all performance against the design targets
- 06 Runbook + documentation handover
- 07 Support tiers

PROCUREMENT OPTIONS

- Outright purchase
- 36-month service agreement
- Design-only engagement if procurement runs elsewhere

