



Intel E835 Private AI Clusters: 200GbE RoCE Guide

Plan and build private AI compute with GPU Smith. We connect workload requirements with GPU hardware, networking, deployment and operating decisions so your infrastructure fits the work it must perform.

Published by GPU Smith · 2026-09-20 · intel e835 private ai clusters · intel e835 · 200gbe roce · ai cluster networking · rdma · rocev2 · private ai infrastructure · gpu cluster networking · sr-io

Original article: <https://gpusmith.com/articles/en/intel-e835-private-ai-clusters>



In this article

1. Executive Summary
 2. Introduction and Background
 3. Product and Platform Architecture
 4. Use Cases and Functional Capabilities
 5. RoCE, iWARP, and Fabric Operations
 6. Virtualization, Lifecycle, and Security
 7. Market Share, Adoption, and Competitive Context
 8. Data Analysis and Evidence
 9. Qualification Pilot and Procurement Decision
 10. Implications and Future Directions
 11. Frequently Asked Questions (FAQs)
 12. Conclusion
-

Executive Summary

Intel launched the **Ethernet E835** controller and adapter family at Computex in **June 2026** (www.intel.com), positioning it from **10GbE through 200GbE** and exposing adapter modes that include **2x25G, 4x25G, 2x100G, and 1x200G** (newsroom.intel.com). The headline rate converts to a theoretical **25 GB/s** before Ethernet, IP, transport, RDMA, and application overhead because an octet contains eight bits (rfc-editor.org). That ceiling is useful for link budgeting, but it is not evidence of model throughput, collective performance, storage goodput, or latency.

The family is best evaluated by traffic role. E835 is a plausible candidate for inference front ends, CPU east-west traffic, storage and model loading, tenant networks, and converged host links. Intel documents concurrent **RoCEv2 and iWARP** plus DPDK-enabled Open vSwitch support (cdrdrv2-public.intel.com) (cdrdrv2.intel.com). It is not yet documented as a drop-in GPU collective fabric. NVIDIA's public GPUDirect documentation explicitly names ConnectX adapters for its peer-memory path, and requires the GPU and peer device to share an upstream PCIe root complex (docs.nvidia.com). No fetched NVIDIA, AMD, Intel, or server-OEM source explicitly qualified an E835 SKU for GPUDirect RDMA, NCCL, RCCL, or a named GPU server topology as of **September 19, 2026**.

The procurement unit is therefore not “an E835.” It is an exact adapter order code, port mode, host slot and NUMA placement, cable or optic, forward error correction mode, switch configuration, operating system, driver, and non-volatile memory (**NVM**) tuple. Intel itself says actual support is determined by NVM and software, not silicon capability alone (www.intel.com). Intel listed an E835-CCQDA1 recommended-customer-price range of **\$553 to \$574**, but a dated channel quote is still required for delivered price, warranty, lead time, and support (www.intel.com). A go decision should follow an acceptance pilot that reproduces the intended server, switch, media, firmware, driver, virtualization, RDMA, and application stack.

Introduction and Background

Private AI clusters carry several networks that are easy to conflate: user and inference traffic, storage reads and checkpoints, cluster management, CPU-to-CPU east-west flows, virtual-machine or container tenant traffic, and GPU collectives. Each path has a different sensitivity to bandwidth, tail latency, loss, congestion, isolation, and operational failure. A NIC described as “AI networking” may fit one path and be unsuitable or simply unproven for another.

This report treats **Intel E835 private AI clusters** as an engineering qualification question, not a branding question. Intel's launch materials describe E835 as supporting Linux, ESXi, and Windows, with a claimed lifecycle longer than ten years (www.intel.com) (www.intel.com). Those are useful fleet-planning signals, but they do not replace an operating-system and release-specific support matrix.

GPU Smith is an adjacent independent engineering advisor, not an E835 vendor. Its published method describes port-by-port fabric validation against a link budget and acceptance testing against written criteria (gpusmith.com) (gpusmith.com). That perspective informs the decision framework here: specify the traffic role first, then prove the complete path.

Product and Platform Architecture

What Intel launched

The E835 generation extends Intel's 800 Series to **200GbE** while retaining the shared driver lineage with E830 and E810 (cdrdv2.intel.com). Intel offers embedded controller and adapter implementations. The two higher-density controller configurations reach **1x200, 2x100, 4x50, or 8x25 GbE** (www.intel.com). Controller capability is not the same as the ports exposed by a retail or OCP card.

Table 1 summarizes the implementation choices that a buyer must pin down before qualification.

Implementation	Physical exposure and intended use	Host and network questions to close
E835-CCQDA1 / E835-CQDA2 , standard or OCP 3.0	One or two high-speed ports, with 200/100/50/25/10GbE listed and DAC reach up to 5 m (cdrdv2-public.intel.com)	Confirm QSFP56 versus QSFP28 implementation, exact order code, supported breakout, switch mode, FEC, optic or cable, slot mechanics, and OEM HCL.
E835-XXVDA4 / E835-XXVDA2 , standard or OCP 3.0	SFP28 implementations for four or two lower-speed ports; OCP XXVDA4 is listed at 25/10/1GbE and four ports (cdrdv2-public.intel.com)	Treat these as tenant, management, or fan-out candidates, not as 200GbE cards. Verify SFP28 module support and port count in the target server.
E835-CCAM1 / CCEM1 / CAM1 , embedded controllers	Silicon for an OEM motherboard or custom adapter, with logical modes set by the implementation	Ask the OEM which controller, SerDes, connector, NVM image, thermal design, port modes, and drivers are actually qualified.

The table shows why family-level procurement is unsafe. For example, Intel's retail CQDA2 page calls its media interface **QSFP28** for direct-attach copper, optics, and active optical cables (www.intel.com). A controller's 8x25 logical capability does not establish that a selected card and cable expose eight independently usable ports.

Host interface and placement

Intel describes E835 host connectivity as PCIe 5.0 or PCIe 4.0 (newsroom.intel.com). PCI-SIG's generation table lists **16 GT/s for PCIe 4.0** and **32 GT/s for PCIe 5.0**, so PCIe 4.0 x16 and PCIe 5.0 x8 have equal aggregate raw signaling before encoding and protocol effects (pcisig.com). That equivalence does not prove equal application behavior.

Qualification must record:

- **Slot wiring:** physical connector, electrical lane width, negotiated generation, and BIOS bifurcation.
- **NUMA locality:** NIC, GPU, storage controller, memory, and CPU socket placement.
- **Root complex:** whether a GPU-direct candidate NIC shares the GPU's upstream PCIe root.
- **Riser limits:** supported card class, airflow direction, power budget, and mechanical clearance.
- **Oversubscription:** simultaneous ingress and egress offered load against measured PCIe, memory, and inter-socket capacity.
- **Isolation settings:** input-output memory management unit mode, access-control services, and address-translation services where virtualization or peer access is used.

Use Cases and Functional Capabilities

The correct port mode follows the traffic pattern. *Table 2* maps common private-AI paths to the evidence required for a decision.

Traffic role	E835 fit hypothesis	Required proof before standardization
Inference/API front end	Strong candidate where 25 to 200GbE Ethernet, RSS, QoS, and predictable tail latency meet the service-level objective	Replay production-sized requests; measure p50, p95, p99, CPU use, drops, retransmits, and failover.
Storage and model loading	Plausible for NVMe-oF, object, NFS, or distributed storage paths using TCP, iWARP, or RoCEv2	Measure warm and cold model loads, checkpoint throughput, CPU cost, queue depth, and congestion interaction with compute traffic.
Cluster control and management	Usually excessive at 200GbE; the 25GbE SFP28 variants may fit consolidated control, provisioning, and telemetry	Test PXE or provisioning, management-plane isolation, upgrade recovery, and out-of-band visibility.
CPU east-west traffic	Good candidate where 1x200 or 2x100 reduces slot count and DDP or RDMA benefits the actual application	Benchmark the buyer's message sizes and concurrency, not a generic packet-rate claim.
GPU collective fabric	Unproven until explicitly validated for the chosen GPU, server, kernel, peer-memory mechanism, and collective library	Compare host-memory and GPU-memory RDMA, then run collective correctness and scaling tests across the intended topology.
Virtualization and tenant traffic	Feature-rich candidate using SR-IOV, VFs, OVS, and DPDK	Validate PF/VF reset, spoof protection, rate limiting, IOMMU containment, noisy-neighbor behavior, telemetry, live operations, and hypervisor support.

Traffic role	E835 fit hypothesis	Required proof before standardization
Precision timing	Candidate where the exact implementation exposes IEEE 1588 and the timing stack is supported	Measure timestamp accuracy, holdover, failover, and switch boundary or transparent-clock behavior.

The important interpretation is that E835 can serve several networks, but [convergence raises the failure radius](#). Intel lists on-chip QoS and traffic management on E835-CCQDA1 (www.intel.com). That feature can enforce a design, but it does not create a correct design. Operators must allocate queues, priorities, bandwidth guarantees, and failure behavior explicitly.

The GPU-direct boundary

Remote direct memory access (**RDMA**) lets one endpoint access another endpoint's memory with reduced CPU involvement. **GPUDirect RDMA** is more specific: the NIC and GPU need a functioning peer-memory path through the PCIe topology and software stack. NVIDIA warns that a path crossing an inter-socket QPI or HyperTransport link may be severely limited or unreliable (docs.nvidia.com). NVIDIA also says IOMMU translation must be disabled or placed in pass-through mode for that documented path (docs.nvidia.com).

Most importantly, NVIDIA's legacy peer-memory documentation explicitly identifies **ConnectX-3 VPI or newer**, not generic RDMA adapters as a class (docs.nvidia.com). Therefore, Intel's RoCEv2 flag cannot be treated as evidence of E835 GPUDirect support. A buyer needs an explicit Intel, GPU-vendor, collective-library, or OEM validation for the exact tuple, or must own the risk through reproducible testing.

RoCE, iWARP, and Fabric Operations

Choose the RDMA mode by failure model

E835 supports concurrent **iWARP and RoCEv2**, but the network dependencies differ. iWARP runs over TCP. In Microsoft's Azure Local design, PFC is optional for iWARP (learn.microsoft.com). RoCEv2 runs over UDP/IP, uses IANA port **4791**, and is routable at layer 3 (www.iana.org) (docs.nvidia.com).

For a RoCEv2 design, specify:

- **Traffic class:** DSCP or priority-code-point mapping from host through every switch hop.
- **PFC scope:** the lossless priority only, never an undocumented blanket setting.
- **ECN behavior:** marking thresholds, NIC reaction, congestion-notification packets, and visibility.
- **Routing:** equal-cost multipath hashing and symmetry for the chosen flow distribution.
- **MTU:** one value across source, switch ports, routed interfaces, and destination.
- **Headroom:** per-priority buffers sized for speed, cable delay, and topology.
- **Failure handling:** pause storms, congestion spreading, mis-marked traffic, and a failed or mismatched switch port.

IEEE defines priority flow control (**PFC**) as intended to eliminate congestion-driven frame loss (1.ieee802.org). It is not risk-free. IEEE working material notes that hop-by-hop backpressure can cause deadlock (www.ieee802.org). Microsoft's implementation guidance requires PFC and Enhanced Transmission Selection

to be consistently configured across hosts and switches (learn.microsoft.com).

MTU, FEC, and media

Intel requires every device in a jumbo-frame path to support and use the same frame size (edc.intel.com). Do not adopt a jumbo MTU by habit. Validate application goodput, fragmentation behavior, and operational tooling at the selected value.

Forward error correction (**FEC**) is equally end-to-end at the link level. Intel says many optical and direct-attach links above 10Gb/s require Reed-Solomon FEC (edc.intel.com). Both link partners must enable a compatible mode. The purchase order should name adapter SKU, cable or optic part number, length, switch SKU, port mode, breakout map, and FEC mode.

Virtualization, Lifecycle, and Security

E835 exposes substantial virtualization capacity. Intel states up to **256 virtual functions**, while its feature matrix records a tested one-port communications case of **64 VMs and 128 VFs**, two VFs per VM (cdrdv2.intel.com) (cdrdv2.intel.com). These are capability limits, not a tenant-density recommendation.

The virtualization pilot should test:

- **PF and VF lifecycle:** creation, bind, unbind, reset, host reboot, and firmware update.
- **Isolation:** MAC and VLAN spoofing controls, DMA containment, IOMMU groups, and rate limits.
- **Queue allocation:** per-VF queues and exhaustion behavior under mixed tenants.
- **Data plane:** kernel, SR-IOV passthrough, DPDK, and OVS paths separately.
- **Observability:** tenant-attributable drops, congestion, errors, and reset events.
- **NUMA pinning:** vCPU, huge pages, poll-mode driver threads, and NIC queues on the local socket.
- **Safety:** avoid DPDK's no-IOMMU mode for multi-tenant use because it removes DMA protection (doc.dpdk.org).

Version control is central. Open vSwitch notes that each DPDK release is validated against a specific firmware version, and warns that a remote-NUMA poll-mode thread reduces maximum throughput (docs.openvswitch.org) (docs.openvswitch.org). Intel likewise says the E835 device image and driver should be updated as a matched set (cdrdv2.intel.com).

As of the publication date, Intel's support catalog listed **ice 2.6.7**, release notes **31.2.2**, and E830/E835 NVM updater **2.11**, all dated July 10, 2026 (www.intel.com) (www.intel.com). A production bill of materials should freeze the exact tuple, not merely "latest."

Security controls include **SPDM 1.2 attestation**, secure boot, secure firmware update, and recovery behavior (newsroom.intel.com) (cdrdv2.intel.com). Intel documents a minimum security revision that cannot be decreased, which affects rollback planning (edc.intel.com). Recovery testing should include a power cycle because Intel requires one after recovery mode (edc.intel.com).

Market Share, Adoption, and Competitive Context

E835 was only launched in June 2026, and the reviewed primary sources disclose no E835 unit shipments, installed-base count, or market share. Public evidence supports product availability and driver release timing, not adoption. Procurement teams should not substitute Intel's long-lifecycle statement for proof of supply in a specific region or OEM platform.

Table 3 places E835 against nearby adapter choices without pretending they are feature-identical.

Family	Official headline interface	GPU-path evidence and procurement interpretation
Intel E835	Up to 200GbE; PCIe 5.0/4.0; retail, OCP, and embedded variants	Strong general Ethernet, RDMA, virtualization, and security feature set. No fetched primary source explicitly qualified E835 for GPUDirect or a named collective stack.
Intel E810-CAM1	Up to 100GbE with 1x100, 2x50, 4x25, or 4x10; PCIe 3.0/4.0 x16 (www.intel.com)	Lower line rate, older platform, shared 800-Series driver lineage. It may remain preferable where the server and software matrix is already qualified.
NVIDIA ConnectX-7	Up to 400Gb/s, four ports, PCIe 5.0 x16 (www.nvidia.com) (networking-docs.nvidia.com)	NVIDIA documents GPUDirect RDMA over RoCE for its ConnectX path (networking-docs.nvidia.com). This is stronger public evidence for NVIDIA GPU collectives, but still needs server and workload qualification.
Broadcom P2200G	Dual-port 200Gb/s or single-port 400Gb/s, PCIe 5.0 x16 (docs.broadcom.com)	A higher-rate alternative on paper. Compare exact offloads, operating systems, optics, OEM HCL, software support, power, and application results rather than headline speed.

The table is not a universal ranking. An already qualified E810 may carry less integration risk than a faster new card. ConnectX-7 has the clearest public NVIDIA GPU-direct documentation among these choices. E835 may still win for a 200GbE storage, inference, CPU, or virtualized network where Intel's driver continuity and port flexibility matter more than a vendor-documented GPU peer path.

Data Analysis and Evidence

The quantitative starting point is line rate: $200\text{ Gb/s} \div 8 = 25\text{ GB/s}$, equivalent to about **23.28 GiB/s** using NIST's binary definition of 1 GiB as 1,073,741,824 bytes (physics.nist.gov). This is an aggregate pre-overhead ceiling. A dual-port card, bidirectional traffic, PCIe read/write asymmetry, NUMA crossing, packet size, and CPU or GPU memory path can each change useful throughput.

Use a simple oversubscription worksheet:

- **Offered ingress:** sum the peak or percentile input demand from every active port.
- **Offered egress:** calculate separately rather than assuming symmetric full duplex.
- **Measured NIC goodput:** payload bytes delivered by the application per second.
- **Measured PCIe path:** negotiated generation and width, plus observed counters and memory bandwidth.
- **Concurrency case:** storage load, inference serving, checkpointing, and tenant traffic active together.
- **Acceptance ratio:** application goodput divided by offered line-rate bytes, explicitly scoped to message size and direction.

Intel's April 2026 test reported **11.68 W incremental loaded power** for E835-CQDA2, but that is a vendor measurement with its own setup, not a universal adapter TDP or buyer benchmark (edc.intel.com). Measure wall or rail power on the target server at idle and under the intended packet mix, then use quoted electricity and support costs. Likewise, Intel's **\$553 to \$574** range is a planning input, not delivered cost.

A useful total-cost worksheet includes:

- **Acquisition:** adapter, riser, cable or optics, switch ports, licenses, spares, and support.
- **Power:** measured incremental server and switch watts, utilization profile, [facility overhead](#), and [local tariff](#).
- **Operations:** design, qualification, monitoring, incident response, firmware maintenance, and change control labor.
- **Capacity:** measured application goodput and tail latency, not nominal bandwidth.
- **Risk:** cost of an unqualified GPU path, unavailable optic, unsupported hypervisor release, or rollback constraint.

No independent E835 workload benchmark or broad adoption dataset was located. That absence is consequential: vendor performance-per-watt or line-rate claims should remain separate from locally measured application results. The same discipline applies to the price range, which excludes the buyer's channel discount, taxes, support, optics, switch capacity, and integration labor.

Qualification Pilot and Procurement Decision

Build the exact bill of materials

Before hardware arrives, freeze and document:

- **Adapter identity:** retail, OCP 3.0, or embedded controller; full order code and hardware revision.
- **Server identity:** exact model, riser, slot, CPU count, NUMA map, BIOS, BMC, airflow, and power policy.
- **Media:** optic, DAC, or AOC part number, reach, breakout, and FEC requirement.
- **Switch:** platform, ASIC, port profile, firmware, VLAN, MTU, QoS, PFC, ECN, and buffer policy.
- **Software:** OS or hypervisor build, kernel, Intel PF/VF/RDMA driver, NVM, DPDK, OVS, RDMA core, GPU driver, and collective library.
- **Workload:** packet and message sizes, concurrency, directionality, latency percentiles, availability target, and failure scenarios.

Dell's general NIC-qualification guidance says to confirm the exact server model in the compatible-systems list (www.dell.com). The reviewed sources did not expose an OEM HCL explicitly naming an E835 SKU, so buyers should obtain written OEM confirmation rather than infer support from PCIe fit.

Execute acceptance tests by role

The pilot should produce raw logs, configurations, and before/after counters for these stages:

1. **Inventory and recovery:** record PCI IDs, serials, NVM, driver, firmware signing state, secure-boot behavior, update, downgrade policy, and recovery procedure.

2. **Physical link:** negotiate every intended port mode and breakout; verify FEC, lane mapping, optic telemetry, temperature, CRC, corrected and uncorrectable blocks.
3. **IP and MTU:** validate end-to-end MTU, routing, equal-cost paths, failover, VLANs, and policy markings.
4. **Baseline transport:** run TCP and UDP throughput, latency, packet-rate, CPU, interrupt, and NUMA tests across representative message sizes.
5. **RDMA:** test iWARP and/or RoCEv2 setup, queue-pair scale, routing, congestion, retries, and error recovery.
6. **Congestion:** create incast, long flows, microbursts, mixed priorities, and a deliberately congested receiver; observe PFC, ECN, CNP, and drops.
7. **Virtualization:** exercise PF/VF isolation, resets, tenant rate control, DPDK/OVS paths, and host maintenance.
8. **GPU path if required:** compare host-memory and GPU-memory RDMA, validate same-root placement, then run [collective correctness and scaling](#).
9. **Resilience:** test link loss, switch reboot, LAG mode if used, driver reload, firmware recovery, and node restart.
10. **Soak:** run the buyer's mixed workload long enough to expose heat, counter drift, memory pressure, and intermittent link errors.

Linux exposes a FEC statistic for blocks that could not be corrected, making it suitable for before-and-after physical-link checks (docs.kernel.org). NVIDIA's troubleshooting guidance recommends comparing host-memory and GPU-memory RDMA results and collecting PFC, ECN, CNP, and queue-drop counters (docs.nvidia.com) (docs.nvidia.com). NVIDIA's nccl-tests checks both collective performance and correctness (github.com).

Go or no-go criteria

Approve E835 for a named role only when:

- **Compatibility is explicit:** every host, adapter, media, switch, driver, NVM, OS, and application component is recorded.
- **Performance is sufficient:** measured throughput, latency, CPU cost, and concurrency meet the written workload target.
- **Congestion is bounded:** loss, pause, ECN, CNP, and recovery behavior remain within operator-defined limits.
- **Isolation is demonstrated:** tenant and management boundaries survive resets, faults, and abusive traffic.
- **Recovery is rehearsed:** update, rollback constraints, failsafe, spare replacement, and power-cycle requirements are operationalized.
- **Support is documented:** OEM, Intel, switch, optics, GPU, and software owners accept the exact configuration.
- **Economics are quoted:** acquisition, support, media, switch capacity, power, and engineering labor use dated inputs.

Reject or defer it for a role when the required GPU peer path is undocumented, an OEM will not support the card in the target server, the intended optic or breakout is absent from compatibility lists, congestion behavior cannot be bounded, or the driver/NVM/hypervisor tuple cannot be frozen and recovered.

Implications and Future Directions

E835 changes Intel's position in host networking by bringing **200GbE** and flexible port modes into the 800-Series lineage. The practical opportunity is consolidation: a private-AI operator may reduce slots or standardize one driver family across inference, storage, CPU, and tenant networks. Any link aggregation design still needs release-specific verification and failure testing.

The near-term uncertainty is ecosystem validation. Public specifications have arrived faster than named OEM server matrices, E835 GPU-direct statements, independent benchmarks, and field-adoption data. Those gaps should narrow as server vendors publish HCLs and operating-system matrices mature. Until then, buyers should preserve a distinction between three statuses:

- **Documented:** the vendor publishes the capability for the exact SKU and release.
- **Interoperable:** the complete physical and software path passes repeatable laboratory tests.
- **Supported:** relevant vendors accept the deployed tuple and its failure modes.

That framework also protects future upgrades. A port profile, optic, or driver change can move a configuration from supported to merely interoperable. Fleet governance should retain golden NVM and driver bundles, configuration exports, telemetry baselines, spare media, and a regression suite for every update.

Frequently Asked Questions (FAQs)

Does Intel E835 support 200GbE RoCE?

Yes, selected E835 implementations reach **200GbE**, and Intel documents concurrent RoCEv2 and iWARP. That does not prove a lossless fabric, GPU-direct path, or application goodput. The exact adapter, port mode, switch, media, FEC, PFC/ECN policy, driver, and NVM must be qualified.

Is E835 compatible with NVIDIA GPU clusters?

It can carry ordinary Ethernet or RDMA traffic in a GPU server if the server and software support it. No fetched primary source explicitly qualified E835 for NVIDIA GPUDirect RDMA or NCCL. NVIDIA's documentation requires topology and peer-memory support, so GPU-memory and collective tests are mandatory.

Should a cluster use RoCEv2 or iWARP?

Choose based on operations. RoCEv2 offers routable UDP/IP RDMA but typically needs a deliberately engineered congestion and QoS policy. iWARP uses TCP and can reduce dependency on PFC. Benchmark both where the application supports them, including failure behavior and CPU cost.

Can E835 replace E810?

Not automatically. E835 doubles the headline controller rate from E810-CAM1's 100GbE to as much as 200GbE and offers newer modes, but an already qualified E810 deployment may have lower integration risk. Compare slot wiring, media, drivers, support, and measured application performance.

What should procurement request?

Request the exact order code, hardware revision, form factor, port profile, media compatibility, FEC mode, power and airflow data, warranty, lifecycle notice, lead time, support entitlement, OS matrix, OEM HCL, and the approved driver/NVM bundle. Require dated quotes rather than relying on a recommended-customer-price range.

Conclusion

Intel E835 is a credible 200GbE candidate for private-AI host, storage, inference, CPU, and virtualized networks, provided the decision is tied to an exact adapter and traffic role. Its flexible port modes, RoCEv2 and iWARP support, SR-IOV scale, OVS/DPDK path, shared 800-Series drivers, timing, and security controls create a broad qualification envelope.

That breadth is also the main procurement trap. Controller capability does not guarantee that a retail, OCP, or OEM implementation exposes the same ports, media, firmware features, or operating-system support. Nor does generic RDMA establish GPU-direct or collective-library compatibility. The absence of public E835 adoption data and independent workload benchmarks makes local evidence more important, not less.

The go/no-go standard is straightforward: freeze the complete hardware and software tuple, reproduce the real traffic mix, measure useful throughput and latency, exercise congestion and failures, prove tenant isolation, rehearse firmware recovery, and obtain written support for the deployed configuration. If GPU collectives are in scope, require explicit vendor qualification or treat host-memory versus GPU-memory RDMA and collective correctness as hard acceptance gates. On that basis, E835 can be standardized for the roles it proves, without mistaking a 200GbE label for a complete AI fabric design.

Source links

Links cited in this article, in order of first appearance. Full addresses are provided for readers using extracted text.

1. <https://www.intel.com/content/www/us/en/newsroom/news/data-center/intel-introduces-ethernet-e835-controllers-network-adapters.html>
2. <https://newsroom.intel.com/data-center/intel-introduces-ethernet-e835-controllers-network-adapters>
3. <https://www.rfc-editor.org/info/rfc1392/>
4. <https://cdrdv2-public.intel.com/709766/Intel%20Ethernet%20800%20Series%20Product%20Guide.pdf>
5. <https://cdrdv2.intel.com/v1/dl/getContent/913514>
6. <https://docs.nvidia.com/cuda/gpudirect-rdma/>
7. <https://www.intel.com/content/www/us/en/content-details/874066/intel-ethernet-controller-e835-feature-support-matrix.html>

8. <https://www.intel.com/content/www/us/en/products/sku/245160/intel-ethernet-network-adapter-e835ccqda1/specifications.html>
9. <https://www.intel.com/content/www/us/en/newsroom/news/data-center/intel-puts-agentic-ai-xeon-6-networking-ai-systems.html>
10. <https://gpusmith.com/>
11. <https://www.intel.com/content/www/us/en/ark/products/series/246043/intel-ethernet-e835-controllers-up-to-200gbe.html>
12. <https://www.intel.com/content/www/us/en/products/sku/245162/intel-ethernet-network-adapter-e835ccqda2/specifications.html>
13. https://pcisig.com/sites/default/files/2026-01/PCI-SIG%20PCIe%207.0%20Webinar_Rev5_FINAL.pdf
14. <https://gpusmith.com/articles/en/gpu-cpu-nic-affinity-rank-binding>
15. <https://gpusmith.com/articles/en/gpu-cluster-reliability-budget>
16. <https://learn.microsoft.com/en-us/azure/azure-local/concepts/host-network-requirements?view=azloc-2605>
17. <https://gpusmith.com/articles/en/infiniband-vs-spectrum-x-vs-roce-vs-ethernet-ai-clusters>
18. <https://www.iana.org/assignments/service-names-port-numbers?search=RoCE>
19. <https://docs.nvidia.com/networking-ethernet-software/cumulus-linux-40/Network-Solutions/RDMA-over-Converged-Ethernet-RoCE/>
20. <https://gpusmith.com/articles/en/roce-fabric-acceptance-testing-runbook>
21. <https://1.ieee802.org/dcb/802-1qbb/>
22. <https://www.ieee802.org/1/files/public/docs2024/dt-seaman-clause-36-proposal-0524-v2.pdf>
23. <https://edc.intel.com/content/www/us/en/design/products/ethernet/adapters-and-devices-user-guide/30.5.1/jumbo-frames/>
24. <https://edc.intel.com/content/www/us/en/design/products/ethernet/adapters-and-devices-user-guide/forward-error-correction-fec-mode/>
25. <https://cdrdv2.intel.com/v1/dl/getContent/874066>
26. https://doc.dpdk.org/guides-21.02/linux_gsg/linux_drivers.html
27. <https://docs.openvswitch.org/en/latest/intro/install/dpdk/>
28. <https://www.intel.com/content/www/us/en/support/products/245454/ethernet-products/800-series-network-adapters-up-to-200gbe/intel-ethernet-e835-network-adapters-up-to-200gbe.html>
29. <https://edc.intel.com/content/www/us/en/design/products/ethernet/adapters-and-devices-user-guide/firmware-security/>
30. <https://edc.intel.com/content/www/us/en/design/products/ethernet/adapters-and-devices-user-guide/intel-ethernet-nvm-update-tool/>
31. <https://www.intel.com/content/www/us/en/products/sku/187409/intel-ethernet-controller-e810cam1/specifications.html>
32. <https://www.nvidia.com/en-gb/networking/ethernet-adapters/>
33. <https://networking-docs.nvidia.com/connectx7hw/introduction>
34. <https://networking-docs.nvidia.com/gpudirectrdma>
35. <https://docs.broadcom.com/doc/PCIe-NIC-Ethernet-Adapters-Specification-Sheet>
36. <https://physics.nist.gov/cuu/Units/binary.html>
37. <https://edc.intel.com/content/www/us/en/products/performance/benchmarks/ethernet/>
38. <https://gpusmith.com/articles/en/gpu-cluster-energy-calculator>
39. <https://www.dell.com/support/kbdoc/en-us/000349619/poweredge-verifying-network-card-compatibility>

- 40. <https://gpusmith.com/articles/en/gpu-cluster-acceptance-testing-runbook>
- 41. <https://docs.kernel.org/6.16/networking/ethtool-netlink.html>
- 42. https://docs.nvidia.com/deeplearning/nccl/user-guide/docs/troubleshooting/networking_troubleshooting.html
- 43. <https://github.com/nvidia/nccl-tests>

THE PUBLISHER

About GPU Smith

GPU Smith provides independent engineering and research for private AI infrastructure. We help technical buyers and operators reason about GPU workload sizing, hardware selection, procurement, deployment and inference operations, from the compute system to the networking, power and cooling requirements around it.

Start with the workload

Useful infrastructure decisions begin with the models, data, throughput, latency and operating constraints a team actually has. GPU Smith helps connect those requirements with system design and deployment choices. The goal is a privately controlled AI environment whose capacity and operational demands are understood before a purchasing decision.

Hardware and supplier research

Our public reference library covers GPUs, complete systems and networking components. The server-vendor directory and manufacturing research help teams investigate suppliers, compare options and follow sources behind technical claims. We distinguish manufacturer specifications from measured benchmarks and advertised capabilities from tested configurations. Reference pages are research resources, not a live inventory listing or a binding equipment quote.

Deployment and operations

GPU Smith's engineering scope includes infrastructure integration, networking, power, cooling and the practical operation of inference workloads. Our published articles explain the tradeoffs behind deployment, procurement and ongoing operations so teams can ask better questions and document decisions.

Work with GPU Smith

Explore the [hardware library](#), [GPU references](#), [networking references](#) and [vendor research](#). [Contact GPU Smith](#) with your workloads and deployment constraints to discuss a project and confirm current scope, pricing and availability.

Inclusion of a manufacturer or product in our research does not imply a partnership, certification or endorsement.

Website: <https://gpusmith.com>

This article is educational. Verify consequential medical, legal, financial, regulatory and technical decisions with appropriate professionals and the cited primary sources. Company services are described separately from the article's research findings.