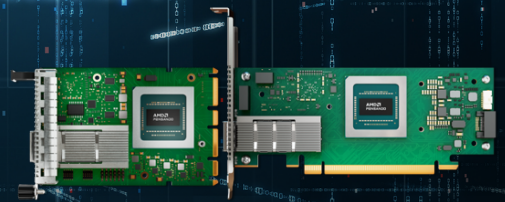


AMD PENSANDO™ POLLARA 400 AI NIC



OVERVIEW

The AMD Pensando™ Pollara 400 AI NIC is a fully programmable up to 400 Gigabit per second (Gbps) Ethernet Network Interface Card (NIC) for scale-out GPU networking and front-end host networking, designed to accelerate network performance in GPU servers for AI data centers.

The AMD Pensando Pollara 400 AI NIC builds on the success of the proven AMD Pensando P4 architecture by combining a high-bandwidth Ethernet controller with a unique set of fully programmable highly optimized hardware acceleration engines.

The AI NIC enhances network performance and helps improve AI job completion times while working as a scale-out AI NIC, and as a Front-End Host NIC, it streamlines connectivity and enables smooth, reliable data movement from the broader network to the AI servers –optimizing overall system efficiency and responsiveness.

Pollara AI NIC is offered in both standard low profile HHHL and OCP-3.0 TSFF form factors.

ADVANCED AI NETWORKING

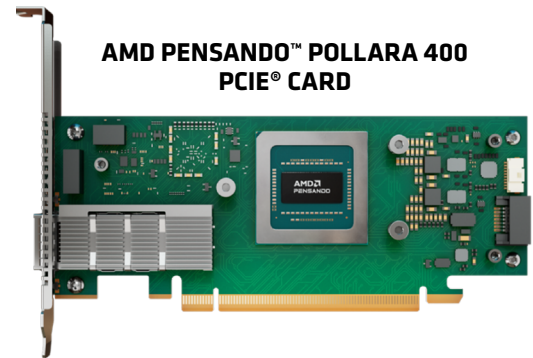
AMD Pensando™ Pollara 400 AI NICs offer an advanced networking solution designed to optimize front-end and back-end networking for GPU-based AI Data Center servers. The AMD AI NIC™ adapters helps facilitate efficient data exchange across scale-out networks through bypassing components not required for GPU-to-GPU communication, enabling reduced latency and increased throughput.

The AMD Pensando Pollara 400 AI NIC is the industry's first Ultra Ethernet Consortium (UEC) ready, fully programmable AI NIC backed by an open ecosystem. Built on the AMD Pensando 3rd generation P4 programmable engine, it provides a foundation for building resilient, scalable and fault tolerant scale-out networks for AI systems by helping to maximize throughput and reduce downtime, which is critical for running uninterrupted training jobs. The AMD Pensando Pollara 400 AI NIC offers unique innovative features, allowing Ethernet to run AI workloads across the back-end infrastructure, including:

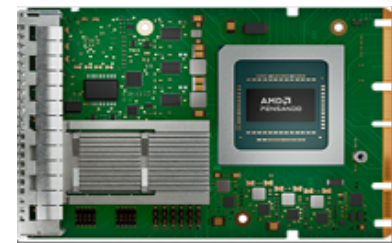
- Intelligent Packet Spray
- In-Order-Delivery (messages to GPU)
- Selective Retransmission
- Path Aware Congestion Avoidance

The AMD Pensando Pollara 400 AI NIC seamlessly integrates into standard GPU servers as both front-end (Host) and back-end (GPU) NIC delivering high-performance networking specifically optimized for AI and ML workloads while reducing complexity at scale. The AMD Pensando Pollara 400 AI NIC offers RoCEv2 compatibility and interoperability with other NICs.

**AMD PENSANDO™ POLLARA 400
PCIE® CARD**



**AMD PENSANDO™ POLLARA 400
OCP FORM FACTOR CARD**



SPECIFICATIONS

MAX BANDWIDTH	• 400 Gbps
FORM FACTOR	• Half-height, half length (HHHL) and OCP-3.0 TSFF
HOST INTERFACE	• PCIe Gen5.0 x16
ETHERNET INTER-FACE	• QSFP112 (NRZ/PAM4 Serdes)
ETHERNET SPEEDS	• 25/50/100/200/400 Gbps
ETHERNET CONFIGURATIONS	• Supports up to 4 ports - 1 x 400G - 2 x 200G - 4 x 100G - 4 x 50G - 4 x 25G
MANAGEMENT	• MCTP over SMBus

KEY BENEFITS

LEADING PERFORMANCE

- Up to **20% better performance** than competitive solutions¹
- **Boost efficiency by up to 25%** with next-generation UEC supported capabilities²

LOWER COST WITH UNMATCHED SCALE

- Up to **50% lower network savings** with multi-plane deployment³
- Scale-out to **20x the GPU Cluster Scale of InfiniBand**⁴

OPERATIONAL RESILIENCE

- Latency metrics, drop statistics
- Extensible API
- Selective acknowledgement and retransmission (SACK)
- Telemetry threshold alerting
- Advanced security features
- **Up to 50% better Reliability and Serviceability (RAS)** in convergence time using UEC-Ready RDMA vs RoCEv2⁵

FEATURES

AMD PENSANDO™ POLLARA 400 AI NIC	
AI NETWORKING	• UEC Ready RDMA • RCCL optimized data path • Path aware and programmable congestion control • Intelligent packet spray • In-Order delivery & Selective Loss Retransmission • AMD AI NIC™ Adapters Direct RDMA
ENHANCED OBSERVABILITY	• Telemetry threshold alerting • Inline Encryption and Decryption • Latency metrics, drop statistics • Rapid Fault Detection
MISC PROTOCOL ACCELERATION	• RDMA • RoCEv2 • Next Gen RDMA and Customized RDMA
CLASSIC NIC FUNCTIONS	• Jumbo Frame Support • Promiscuous mode • PFC and link-level pause • IP, TCP, UDP, MAC, VLAN filtering • Interrupt coalescing • Receive side scaling (RSS) • Multicast and All-multicast • Multi queue support • Hardware timestamping • Scatter-gather
VIRTUALIZATION	• SR-IOV • Support for multiple profiles of PF, VFs and queues.
HOSTNIC OFFLOADS	• TCP and UDP Checksum offloads (CSOs) • TCP Segmentation offload (TSO/LSO) • Generic receive offloads (GRO) and Large Receive offload (LRO) • VLAN offload and VLAN tunnel for VFs • Stateless tunnel offload (VXLAN, GRE, NVGRE, IP-IP, Geneve)
SOFTWARE AND DRIVERS	• DPDK, XDP, AF-XDP • Virtio-Net
STORAGE NETWORKING	• RDMA and RoCEv2 • Direct RDMA • Direct Storage

MANAGEMENT AND CONTROL

- MCTP/SMBus
- SPDM over MCTP
- MCTP over PCIe VDM
- PLDM firmware
- PXE and IPXE boot

DEPLOYMENT OPTIONS

AI COMPUTE NODE SCALE-OUT (BACK-END) NIC	AI COMPUTE NODE HOST (FRONT-END) NIC
The scale-out AI NIC can be installed for every GPU in a multi-GPU AI node. The AMD Pensando Pollara 400 AI NIC supports remote direct-memory access (RDMA) over Ethernet delivering optimal AI workload efficiency. The AMD Pensando Pollara 400 AI NIC software interacts with collective communication libraries to optimize job completion times at very low latency	To increase time to market and increase ease of use for deploying the same hardware and software, the AMD Pensando™ Pollara 400 AI NIC can also be deployed at front-end networks enabling up-to 400G CPU communication with host NIC functionality which includes services (network, security and storage) offloads in bare-metal and virtualized environments.

POLLARA AI NIC ORDERABLE PART NUMBER (OPN)	
POLLARA-400-1Q400P	• All adapters are shipped with the tall bracket mounted and a short bracket as an accessory

ENDNOTES

- PEN-013: The AMD Pensando™ Pollara 400 delivers substantial cluster communication performance improvements versus Broadcom Thor2, achieving 1.2x the performance (20% higher performance) of Thor2 on a 128 GPU cluster with standard RoCEv2 workloads for most collective communication operations.

Net benefits for customers are fast data transfers, accelerated job completion times and enhanced scalability for high-performing AI workloads with RoCEv2 today, while being future-ready for UEC.
- PEN-016: Testing conducted by AMD Performance Labs as of [28th April 2025] on the [AMD Pensando™ Pollara 400 AI NIC], on a production system comprising of: 2 Nodes of 8xMI300X AMD GPUs (16 GPUs): Broadcom Tomahawk-4 based leaf switch (64x400G) from MICAS network; CLOS Topology; AMD Pensando Pollara AI NIC - 16 NICs; CPU Model in each of the 2 nodes - Dual socket 5th gen Intel® Xeon® 8568 - 48 core CPU with PCIe® Gen-5 BIOS version 1.3.6 ; Mitigation - Off (default)

System profile setting - Performance (default) SMT- enabled (default); Operating System Ubuntu 22.04.5 LTS, Kernel 5.15.0-139-generic.

Following operation were measured: Allreduce

Average 25% for All-Reduce operations with 4QP and using UEC ready RDMA vs the RoCEv2 for multiple different message size samples (512MB, 1GB, 2GB, 4GB, 8GB, 16GB).
The results are based on the average at least 8 test runs.
- PEN-018: AMD comparison and pricing as of July 6, 2025, for network fabric costs to support 128,000 GPUs. Comparison of a Pollara NIC with multiplane fabric and packet spray on an 800G Tomahawk 5-based multiplane design versus a generic fat-tree fabric built on fully scheduled, big-buffer (Jericho3/Ramon3) 800G switching platforms. The generic system is assumed to use a competitive NIC, with NIC costs considered comparable. The Pollara-based design is estimated to deliver up to 58% network switching cost savings by enabling the use of more cost-effective Tomahawk 5-based switching in a multiplane architecture. .AMD comparison and pricing as of 4/23/2025 of a Tomahawk 5 system with Pensando Pollara NIC featuring exclusive multiplane fabric and packet spray versus a generic big-buffer 800G switching platform; the generic system would employ a competitive NIC, costs of NICs are assumed to be comparable. Deploying Pollara with multi-fabric support and packet spray, allows customers to build cost-effective multiplane network fabrics, instead of a fat-tree design using less network switches to deliver the same amount of network bandwidth across the fabric, and dramatically reducing both switch platform cost, and cost associated with cables, optics.
 - Fat-Tree Big Buffer Fully Scheduled Network (Leaf/Spine/Core) estimated Cost: \$1.22B
 - 3556 leaf (Jericho3-AI) units at \$104,998 each = \$373M
 - 1557 spine/core (Ramon3) units at \$147,998 each = \$247M
 - 128K AOC-10m cables at \$1059 each = \$136M
 - 568,889 QDD-SR4-400G transceivers at \$819 each = \$466M
 - Total (Switching & Optics) = \$1.22B
 - Naddod Tomahawk5 800G Multiplane Fabric Network estimated Cost: \$511M
 - 3,000 Leaf and Spine Units (Naddod N9600-640C) at \$26,999 each = \$81M
 - 384K (QDD-SR4-400G) transceivers at \$819 each=\$313M
 - 64K Switch (QSFP-2x400G-DR4) transceivers for NIC connections at \$759 each = \$48M
 - 256K MPO Cables at \$26 each = \$6.6M
 - 2K Optical Shuffle box, modules and internal cables at \$30K per rack =\$60M
 - Total (Switching & Optics) = \$511MPrices subject to change. Comparison for specific network configurations only, and may not be representative of all possible network configurations and comparisons.
- InfiniBand Scale Size 48,000 nodes:
https://network.nvidia.com/pdf/whitepapers/InfiniBandFAQ_FQ_100.pdf
100K RoCEv2: <https://www.tomshardware.com/desktops/servers/first-in-depth-look-at-elon-musk-100-000-gpu-ai-cluster-xai-colossus-reveals-its-secrets>
UEC Scales to 1,000,000 nodes: https://network.nvidia.com/pdf/whitepapers/InfiniBandFAQ_FQ_100.pdf
- PEN-019: Testing conducted by AMD Performance Labs as of [15 September 2025] on the AMD Pensando Pollara AI NIC, on a test system comprising of SMC-300X server for GPU-GPU communication: 2x AMD Pensando Pollara AI NIC, 2P AMD EPYC 9454 48-Core -2P Processor, 8x AMD Instinct MI300X GPU, Ubuntu 22.04.5 LTS, kernel 5.15.0-139-generic, ROCm 6.4.1.0-83-69b59e5. Testing running Llama-3.1-8B, Model Configuration: SEQ_LEN=2048, TP=1, PP=1, CP=1,FP8=1, MBS=10, GBS = 5120. Iteration = 2, No. of paths/QP : 128. Results may vary based on factors including but not limited to system configuration and software settings.

DISCLAIMERS

The information presented in this document is for informational purposes only and may contain technical inaccuracies, omissions, and typographical errors. The information contained herein is subject to change and may be rendered inaccurate for many reasons, including but not limited to product and roadmap changes, component and motherboard version changes, new model and/or product releases, product differences between differing manufacturers, software changes, BIOS flashes, firmware upgrades, or the like. Any computer system has risks of security vulnerabilities that cannot be completely prevented or mitigated. AMD assumes no obligation to update or otherwise correct or revise this information. However, AMD reserves the right to revise this information and to make changes from time to time to the content hereof without obligation of AMD to notify any person of such revisions or changes.

THIS INFORMATION IS PROVIDED "AS IS." AMD MAKES NO REPRESENTATIONS OR WARRANTIES WITH RESPECT TO THE CONTENTS HEREOF AND ASSUMES NO RESPONSIBILITY FOR ANY INACCURACIES, ERRORS, OR OMISSIONS THAT MAY APPEAR IN THIS INFORMATION. AMD SPECIFICALLY DISCLAIMS ANY IMPLIED WARRANTIES OF NON-INFRINGEMENT, MERCHANTABILITY, OR FITNESS FOR ANY PARTICULAR PURPOSE. IN NO EVENT WILL AMD BE LIABLE TO ANY PERSON FOR ANY RELIANCE, DIRECT, INDIRECT, SPECIAL, OR OTHER CONSEQUENTIAL DAMAGES ARISING FROM THE USE OF ANY INFORMATION CONTAINED HEREIN, EVEN IF AMD IS EXPRESSLY ADVISED OF THE POSSIBILITY OF SUCH DAMAGES.

COPYRIGHT NOTICE

AMD, the AMD Arrow logo, Pensando and combinations thereof are trademarks of Advanced Micro Devices, Inc. PCIe® is a trademark of PCI-SIG Corporation. Other product names used in this publication are for identification purposes only and may be trademarks of their respective companies. Certain AMD technologies may require third-part enablement or activation. Supported features may vary by operating system. Please confirm with the system manufacturer for specific features. No technology or product can be completely secure. © 2025 Advanced Micro Devices, Inc. All Rights Reserved.