



Dawning AI Server Network Card Configuration

Article Overview

For optimal AI performance, Dawning AI servers should use high-bandwidth, low-latency NICs with RDMA/RoCEv2 enabled, mapped per GPU via PCIe switches, and integrated into a lossless scale-out fabric.

Network Fabric Overview

Dawning AI servers rely on multiple network fabrics to support AI workloads:

- o Front-End Fabric: Connects applications, storage, and GPU cluster management. Typically requires 1-2 NICs per server .
- o Internal-AI Fabric: Integrates GPUs, CPUs, NICs, and storage via PCIe switches for efficient peer-to-peer communication .
- o Scale-Up Fabric: Handles intra-server GPU-to-GPU communication using NVLink, Infinity Fabric, or Ethernet .
- o Scale-Out Fabric: Connects multiple AI servers for distributed workloads. Requires one NIC per GPU with matching PCIe bandwidth (e.g., PCIe Gen5 x16 or 400Gbps) to ensure non-blocking, low-latency transfers .

NIC Configuration Guidelines

- o RDMA and RoCEv2:
 - o Enable RoCEv2 on Broadcom or compatible NICs to offload transport tasks from the CPU and provide direct memory access between GPUs and network .
 - o Ensure NICs support UDP/IP routable RoCEv2 for cluster-wide communication.
- o PCIe Mapping:
 - o Each GPU should have a dedicated NIC connected through a PCIe switch to avoid bottlenecks .
 - o Verify PCIe bandwidth matches NIC throughput (e.g., 400Gbps NICs require PCIe Gen5 x16 lanes).
- o Server BIOS Settings:
 - o For AMD-based servers (e.g., Supermicro AS-8125GS-TNMR2), disable IOMMU and ACS if required by the workload .
 - o Dell XE9680 servers may require careful SR-IOV configuration; consult vendor documentation to avoid boot faults .
- o GPU-NIC Topology:
 - o Use tools like rocm-smi for AMD GPUs to verify hardware topology and NIC mapping .
 - o Ensure GPUs and NICs are correctly paired for optimal RDMA performance.



Dawning AI Server Network Card Configuration

Network Design Considerations

- o Lossless Ethernet: Configure Priority Flow Control (PFC) and Explicit Congestion Notification (ECN) on switches to prevent packet loss during high-throughput AI workloads .
- o Spine-Leaf Architecture: Recommended for scale-out fabrics to reduce latency and provide sufficient bandwidth for multi-node training .
- o Quality of Service (QoS): Assign DSCP values (e.g., CS3) for RoCEv2 traffic and configure class-maps and policy-maps on switches to maintain low-latency, high-priority traffic .

Monitoring and Validation

- o Use NIC and GPU management tools to monitor throughput, latency, and RDMA performance.
- o Validate network paths and congestion management to ensure non-blocking communication across the AI cluster . By following these guidelines, Dawning AI servers can achieve high-throughput, low-latency networking, essential for distributed AI training and inference workloads.

But as demand for infrastructure capable of supporting AI model training and inference has grown, the ability to host

For optimal performance, the server's configuration should include a balanced PCIe topology,

This guide provides instructions for configuring network card settings to optimize network performance and connectivity.

Learn how to setup and optimize GPU servers for AI integration. This guide covers hardware selection, OS & drivers installation, AI

This article explores the hardware topology and cluster networking of high-performance GPU servers, focusing on

Hygon's expertise in advanced chip design complements Dawning's server and storage solutions, creating a

#41 Hardware requirements for Client AI can vary significantly dependent on the type, size and complexity of project

Transforming a list of carefully selected components into a functional server requires a methodical assembly and configuration



Dawning AI Server Network Card Configuration

Join DAWN, the decentralized wireless network revolutionizing Internet service. Empower communities to buy, sell, and share

BMC network access is provided through the out-of-band network (marked in purple). The networking ports and their

Learn how to set up and optimize GPU servers for AI. Discover best practices for configuration, tuning, and deployment.

A host of power and networking semiconductors enable efficient energy delivery and seamless interconnection within AI servers and

Networking Purpose-Built for the Era of AI Factories AI factories now span tens of thousands of GPUs--and soon,

Learn how to build a high performance AI server to allow you to run large language models locally. Removing the need

This reference design leverages DriveNets AI Fabric and the breakout capabilities of NCP5-AI leaf switches to create a

In some scenarios, you may require two network cards for your setup. This article describes how to configure dual

Web: <https://supremeacademicresearch.co.za>

