

AI Core Computing Server

Ordering information

Model	SP1280	SP2560	SP5120	SP10240	SP20480	SP40960
Product name	Patch Panel	Patch Panel	Patch Panel	Patch Panel	Patch Panel	Patch Panel
Illustration						
HU	1	2	4	1	2	4
Maximum number of cores	144	288	576	144	288	576
Product size (including modules and aspers)	482.6*100*74 mm	482.6*100*183 mm	482.6*100*177 mm	482.6*100*74 mm	482.6*100*183 mm	482.6*100*177 mm
Standard color code	RAL9005	RAL9005	RAL9005	RAL9005	RAL9005	RAL9005
Inventory	/	/	/	/	/	/

AI Overview

An AI Core Computing Server is a specialized, high-performance system designed to handle intensive AI workloads using GPUs, high-speed memory, and advanced interconnects.

Overview AI Core Computing Servers are purpose-built for artificial intelligence tasks, including training large language models, generative AI, deep learning, and real-time inference. Unlike traditional servers that rely primarily on CPUs for sequential processing, AI servers leverage massively parallel processing through GPUs, AI accelerators, and sometimes FPGAs to efficiently handle complex computations and large datasets.

Key Components

- CPUs and GPUs:** High-core-count CPUs manage general tasks, while multiple GPUs (e.g., NVIDIA RTX Pro 4500/6000 Blackwell or Grace Blackwell architectures) handle parallel AI computations.
- Memory:** Large amounts of RAM, often 128GB or more, and high-speed memory interfaces like LPDDR5X, support rapid data access for AI models.
- Storage:** Ultra-fast NVMe SSDs (up to 4TB or more) provide low-latency access to massive datasets.
- Networking and Interconnects:** High-bandwidth, low-latency fabrics such as PCIe Gen 5, NVLink, and RDMA over AI fabric enable fast communication between GPUs and CPUs, crucial for distributed AI workloads.
- Software Stack:** AI servers include optimized drivers, frameworks, and monitoring tools to manage workloads efficiently and support AI-specific operations.

Architecture and Scalability AI servers are often deployed in clusters, allowing multiple nodes to work together as a cohesive system. Modular designs, like the Cisco UCS X-Series, allow mixing CPU and GPU nodes in a single chassis with dynamic resource allocation via X-Fabric interconnects. This architecture supports scalable AI training, inference, and visualization workloads without re-architecting the infrastructure.

Use Cases

- Large Language Model Training:** Parallel GPU processing accelerates model training for billions of parameters.
- Generative AI and Deep Learning:** High-density GPU servers enable real-time gener...

Article Content

Document Library

Last Updated: July 07, 2026 The BMX-P820 is built to deliver high-performance computing with efficiency, reliability, and flexibility for demanding industrial AI workloads. It supports 14th Gen

Arm Launches First Server CPU, Expanding Its Push

The move puts Arm deeper into the AI data center market and raises new questions about competition, customers and its relationships with existing

Processors

Built for performance. Ampere's single-threaded, highly parallel cores are designed to deliver consistently fast, low-latency processing from cloud to edge. Designed for AI Compute means

L4 Tensor Core GPU for AI & Graphics | NVIDIA

Accelerate Video, AI, and Graphics Workloads The NVIDIA L4 Tensor Core GPU powered by the NVIDIA Ada Lovelace architecture delivers universal, energy-efficient acceleration for video, AI,

PyImageSearch

Deep Learning Deep Learning algorithms are revolutionizing the Computer Vision field, capable of obtaining unprecedented accuracy in Computer Vision tasks, including Image Classification, Object

AI Computing Servers: The Definitive Guide to Powering Next-Gen AI

With over 8 years of experience in enterprise server solutions, we specialize in providing high-quality, original servers, storage, switches, GPUs, SSDs, HDDs, CPUs, and other IT hardware to clients

NVIDIA Grace CPU and Arm Architecture | NVIDIA

NVIDIA Grace uses high-performance, Arm-based CPU cores and NVIDIA SCF to deliver efficient performance for accelerated computing, cloud, edge, and data center workloads.

AI Server

Altos offers a range of powerful and flexible AI server solution, designed to meet the demands of high-performance computing. Supporting the latest CPU and GPU

PowerEdge AI Servers with GPU Acceleration | Dell USA

Drive faster results with servers equipped with the latest multi-core processors. Leverage CPUs to deliver rapid data processing and enhanced application performance for trusted artificial intelligence

China's dual 16-core Hygon CPU server rack barely outperforms a

A rack comprised of two Hygon server-based processors was spotted on the Geekbench browser featuring a run of Geekbench's new AI-based benchmark.

Cloud and AI Development Act | Shaping Europe's digital future

The Cloud and AI Development Act will strengthen Europe's sovereignty and competitiveness in the cloud and AI ecosystem.

headTitleNoCommunity

Learn about 77 new offers that went live in Microsoft Marketplace, a single destination to find, try, and buy cloud solutions, AI apps, and agents to meet...

NVIDIA Blackwell Platform Arrives to Power a New Era

Powering a new era of computing, NVIDIA today announced that the NVIDIA Blackwell platform has arrived — enabling organizations everywhere to

TechTarget

TechTarget provides purchase intent insight-powered solutions to identify, influence, and engage active buyers in the tech market.

Trump's Remark Ignites Dell Shares: AI Server Leader Surges 13% to ...

TradingKey - On May 8, Eastern Time, U.S. President Trump's remark to "go buy a Dell (DELL) computer" caused shares of the core AI server supplier to jump more than 14% intraday.

AMD EPYC™ 9004 Server CPUs for AI and Cloud Performance

AMD EPYC 9004 Server CPUs deliver high performance, energy efficiency, and low TCO for AI, cloud, and enterprise workloads with up to 128 "Zen 4" cores.

Red Hat OpenShift | comprehensive application platform

Red Hat OpenShift is an application innovation platform delivering a comprehensive experience with tools to empower organizations to modernize and build new

AI infrastructure compute strategy | Deloitte Insights

Instead, it's building infrastructure that leverages the right compute platform for each workload. While exploring AI-optimized infrastructure,

AI Servers with NVIDIA HGX, OAM & PCIe GPU

Explore GIGABYTE AI servers for AI training and inference, supporting NVIDIA HGX, OAM, and PCIe GPU platforms. Compare models to find the right

News Archive

NVIDIA BioNeMo Agent Toolkit Brings Accelerated AI to Life Sciences Researchers in Claude Science Life sciences has entered an era of

Efficient Log Management with Microsoft Fabric

Introduction In the era of digital transformation, managing and analyzing log files in real-time is essential for maintaining application health,...

Intel previews Computex 2026 lineup across handhelds, desktops

Intel will show Panther Lake Arc G3 handhelds, a 52-core Nova Lake desktop, and 288-core Clearwater Forest Xeon at Computex 2026. All are built on the 18A node.

What is an AI Server? AI Server Architecture Explained

In this quick guide, we'll walk you through everything you need to know before deploying your first AI server configuration, covering most of your

Building the AI Server

AI/ML demands are reshaping servers. Explore how CPUs, GPUs, FPGAs and AI accelerators drive performance for workloads like deep learning and predictive analytics.

Contact Us

For more information, pricing, or custom solutions, please contact us:

Website: <https://adikor.eu>

Email: info@adikor.eu

Phone: +33 6 48 37 92 15

Address: 12 Rue de la Paix, 75002 Paris, France

This document is for informational purposes only. Specifications subject to change without notice.

