

Getting Started with AI Servers



Article Overview

To start using an AI server, you can set up a local server using tools like Ollama or build a high-performance server tailored to your needs.

Setting Up a Local AI Server

- **Choose Your Software:** A popular choice for beginners is Ollama, which allows you to run AI models locally without needing extensive technical knowledge. You can download it from ollama.com and follow the installation instructions for your operating system .
- **Install Docker:** If you want to enhance your setup with a graphical interface, consider using OpenWebUI. This requires Docker, which simplifies running applications. Make sure Docker Desktop is running, and follow the installation instructions on docs.openwebui.com to set it up .
- **Run Your AI Models:** After installation, you can start using AI models. For example, you can add the Llama 3.2 model and begin interacting with it through your terminal. This setup allows you to ask questions and receive responses directly from your local server .

Building a High-Performance AI Server

If you're looking for a more robust solution, consider building a high-performance AI server. Here are some key components to consider:

- **Hardware:** Use a powerful CPU, such as the AMD Ryzen 9 7950X, and consider adding dual NVIDIA RTX 4090 GPUs for handling demanding AI workloads. Aim for at least 128GB of RAM to ensure smooth multitasking

Article Content

Local AI Server A Step by Step Guide to Setup and Use

Learn to set up and use your local AI server with this comprehensive guide. Enhance your projects today—read the

Rent GPUs | Vast.ai

Rent high-performance cloud GPUs at low cost with Vast.ai. Instantly deploy GPU rentals for AI, machine learning, deep learning,

Getting started | Red Hat AI Inference Server | 3.0 | Red Hat

Red Hat AI Inference Server is a container image that optimizes serving and inferencing with LLMs. Using AI Inference Server, you

Tools for AI Workstations | NVIDIA

Agentic RAG for improved response quality: Use AI agents for dynamic information retrieval across documents and web search. Run

Unity AI Open Beta: How to get started with MCP

In today's article on the Unity AI Open Beta, learn how to connect Claude Code, Github Copilot, and other AI agents

Getting started with AI in SQL Server 2025 on Windows

It's a guide for setting up SQL Server 2025 in a Hyper-V VM (I used Windows Server) with Ollama and nginx as an

Get Started With NVIDIA AI Enterprise

Get started for free with NVIDIA-hosted NIM™ microservices for accelerated AI models and NVIDIA Blueprints for sample AI

Getting Started with MCP Servers: Build Your First AI Agent in 30 ...

The Model Context Protocol (MCP) is revolutionizing how we interact with AI assistants. Here's your practical guide to

Get started with AI and MCP

Learn about AI and MCP key concepts and development resources to get started building MCP clients and servers

Get started with the remote Power BI MCP server

Learn how to set up and configure the remote Power BI MCP server to enable AI agents to query Power BI semantic

Home AI Server Build Guide 2026 — Always-On Local LLM

Build a dedicated home AI server that runs 24/7 — serving LLMs to every device on your network. Hardware picks,

Getting started with ChatGPT

Learn how to use ChatGPT, start your first conversation, and discover simple ways to write, brainstorm, and solve

Getting started | Red Hat AI Inference Server | 3.2 | Red Hat

Abstract Learn how to work with Red Hat AI Inference Server for model serving and inferencing.

Foundry Agent Service | Microsoft Azure

Use Foundry Agent Service to design, deploy, and scale AI agents securely. Create enterprise-grade AI agents to automate complex

How to Build an Affordable Custom AI Server for AI Projects

In this overview, Jun Yamog guides you through the essentials of building a high-performance AI server, from selecting

How to Build an Affordable Custom AI Server for AI Projects

Take control of your AI projects with a custom-built server. Learn to optimize hardware, reduce costs, and future-proof

Server with GPU: for your AI and machine learning

Get AI models and tools such as DeepSeek or Ollama running on our dedicated GPU servers and tag us

Self-Hosted AI Models: A Practical Guide to Running LLMs Locally

Learn how to self-host AI models for better data control and lower costs. Covers hardware requirements, open-source

Get Started with AI Architecture Design

Get started with AI architecture design on Azure. Explore AI services, reference architectures, best practices, readiness

Tools for developers to get started — Google AI

Discover tools and resources to build with Google AI, customize models, and leverage the power of artificial intelligence.

Full Tutorial

In this video, I'll show you step-by-step how to build your own AI server at home for less

Getting started | Red Hat AI Inference Server | 3.0 | Red Hat

The following troubleshooting information for Red Hat AI Inference Server 3.0 describes common problems related to model loading,

Getting Started with Lemonade | AMD AI Playbooks

Getting Started with Lemonade Learn to run Gen AI models locally with Lemonade, an open-source local AI server.

Getting Started with Google MCP Servers

In this codelab, you will learn how to supercharge your AI agents using Google Managed Model Context Protocol

Contact Us

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

Website: <https://sdesign.net.pl>

Email: sales@sdesign.net.pl

Phone: +48 694 237 815

Address: ul. Marszałkowska 10, 00-001 Warsaw, Poland

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

