

Local Deployment of AI on Cloud Servers



Overview

This article covers the AI workloads available on Azure Local and helps you pick the right one for your needs. Local server deployment for AI means installing and operating GPU servers, storage, and networking at an organization's own facility rather than in a remote or hosted data center. For enterprise teams, local deployment offers direct physical control, ultra-low latency, and the ability to keep. This guide covers what it takes to self-host AI models: the tools, the infrastructure, the costs, and the honest trade-offs you need to consider before making the switch. Every prompt you send to ChatGPT, Claude, or Gemini. Azure Local brings Azure AI capabilities directly to your infrastructure so you can process data locally without sending it to the cloud. Your data is sent to the cloud where powerful data center resources process it, and results are returned over the internet. AI models have advanced extraordinarily quickly — driven by improvements in GPU server technology, the availability of large pre-trained models, and the. As AI integration deepens in late 2025 and early 2026, the discussion around where AI models operate—locally on your device or remotely in the cloud—has intensified.



Article Content

Modal: High-performance AI infrastructure

AI infrastructure that developers love Run inference, training, batch processing, and sandboxes with sub-second cold starts, instant

AI Workloads on Azure Local Overview | Microsoft Learn

Learn how Azure Local enables AI inference, agentic retrieval, and video analysis on your own infrastructure with cloud

Local AI Hosting: How To Run AI Models Yourself

What Is Locally Hosted AI (and Why You Should Care) Locally hosted AI means running machine learning models

AI Workloads on Azure Local Overview | Microsoft Learn

Azure Local brings Azure AI capabilities directly to your infrastructure so you can process data locally without sending it

Private LLM Deployment: Enterprise Self-Hosted AI (2026)

Private LLM deployment for enterprise: run self-hosted AI on-premise for HIPAA, CMMC & GDPR compliance.

Local AI vs Cloud AI: Privacy, Speed, Cost 2026

As AI integration deepens in late 2025 and early 2026, the discussion around where AI models operate—locally on

Tutorial: Build a Low-Cost Local LLM Server to Run 70B Models

Learn how to repurpose crypto-mining hardware and other low-cost components to build a home server capable of

NVIDIA Partners With Microsoft on Unified Stack for Agentic AI ...

NVIDIA Partners With Microsoft on Unified Stack for Agentic AI Deployment, From Windows Devices to Cloud to Local

Cloud AI vs Local AI: Latency, Performance, and Business Impact

Your AI. Your Data. In today's competitive landscape, enterprise leaders face a critical decision when implementing

Cloud AI vs Local AI: Latency, Performance, and

In this deep dive, we explore why organizations are bringing AI in-house, how they're

Edge AI vs. Cloud AI | IBM

Edge AI and cloud AI are two types of AI deployments that have become critical to the development of most modern AI applications.

How to Deploy AI Models: From Training to Production | Dokploy

Learn how to deploy AI models locally, in the cloud, or on your own servers with a practical guide to packaging,

From Local Dev to Production: How to Deploy AI

AI models are more accessible than ever, but taking one from your local machine to

Deploying AI Models on GPU Servers: Step-by-Step Guide 2026

In 2026, the standard for deploying AI models at scale is GPU server infrastructure — and this guide walks you

Run AI Models Locally: Setup & Tips (2026)

Running AI models locally means deploying and operating AI directly on infrastructure the organization owns or fully

Cloud AI vs Local AI: Privacy, Compliance, and Cost

In this deep dive, we'll compare cloud and local AI from a business perspective, focusing on

Local AI Server for Business 2026 — Build Guide + ROI

Build a \$2,500-\$10,000 local AI server for 5-50 users. RTX 4090/5090 hardware tiers, Ollama + Open WebUI stack,

Choose between cloud-based and local AI models

Guidance to help developers that want help choosing between cloud-based and local AI models for their Windows

The Complete Developer's Guide to Running LLMs Locally

A comprehensive guide covering the local LLM stack from hardware requirements to production deployment.

Cloud vs. Local AI: How to Choose Your Deployment Strategy

A practical framework for executives choosing between cloud AI APIs and local model deployment. Covers data

Microsoft Sovereign Private Cloud scales to thousands of nodes with ...

Today, I am pleased to announce that Azure Local now scales to support deployments of up to thousands of servers

Deployment of Machine Learning Models: On-Premises

Whether you choose on-premises or cloud-based deployment, understanding the process

Running LLMs Locally in 2026: Ollama, llama.cpp, and

Run LLMs on local hardware for privacy, lower costs, and faster inference—this guide

Running LLMs on Your Own Hardware: What Actually Works in 2026

A practical guide to running LLMs and AI models locally on your own hardware. Covers Ollama, LM Studio, llama.cpp,

Azure Local Disconnected Operations: Running

Many organizations operate in environments where continuous cloud connectivity is not

How to run LLMs locally: Hardware, tools and best practices

Learn how to run a large language model (LLM) locally, including GPU requirements, multiuser scaling tools and

How to Deploy a Machine Learning Model for Free – 7 ML Model Deployment ...

How to Go from Zero to Hero with Google Cloud Platform How to Deploy Fast.ai models to Google Cloud Functions

Artificial Intelligence Series (Chapter 0): Local

Protecting Sensitive Data: When dealing with sensitive data (such as healthcare, finance, or

How To Deploy a Local AI via Docker

Learn to deploy your own local AI service using Docker containers for maximum security and control, whether you're

What Is Azure Local? Overview and Key Benefits

Learn how Azure Local accelerates cloud and AI innovation by delivering applications, workloads, and services from

Cloud vs. Local: The GenAI Architecture Dilemma | ITSG Global

Cloud vs. Local: The GenAI Architecture Dilemma As companies rush to integrate Generative AI (GenAI) into their

LocalAI

The open, modular AI runtime. Run text, vision, voice, image, video, agents, and more on hardware you control.

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

Your applications call the same endpoints, just pointed at your local server instead of cloud services. For a complete

What is edge AI?

Edge AI refers to the deployment of AI models directly on local edge devices to enable real

Self-Hosted AI: A Complete Roadmap for Beginners

Image by Author # Introduction Building your own local AI hub gives you the freedom to automate tasks, process private data, and

LLM On-Premise : Deploy AI Locally with Full Control

Explore deploying LLMs on-premise: a comprehensive guide to setting up large language

Local Server Deployment for Enterprise AI: Key Requirements

This article covers when local server deployment is the right choice for AI workloads, what facility and infrastructure

Contact Us

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

Website: <https://benniedesignstudio.co.za>

Email: sales@benniedesignstudio.co.za

Phone: +27 82 391 4765

Address: The Towers, 1 St Georges Mall, Cape Town, 8001, South Africa

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

