

HPE ProLiant DL380 Gen10 Server Overview

The **HPE ProLiant DL380 Gen10** is a high-performance computing server designed to handle demanding workloads, including complex machine learning (ML) algorithms, data analytics, virtualization, and enterprise applications. Its robust configuration and scalability make it ideal for both research and production environments.

Specifications

1. Processor

○ 2 x Intel Xeon Silver Processors

- Multi-core processors designed for high-efficiency and parallel processing.
- Suitable for compute-intensive applications like ML model training, big data analytics, and scientific simulations.
- Advanced features include support for Intel Turbo Boost, Hyper-Threading, and virtualization.

2. GPU

○ NVIDIA Tesla GPU (16 GB)

- Specifically designed for high-performance computing (HPC) and machine learning workloads.
- Supports CUDA and Tensor Core acceleration for faster training and inference of deep learning models.
- Optimized for frameworks like TensorFlow, PyTorch, and MATLAB.

3. Memory

○ 128 GB DDR RAM

- High-capacity memory ensures smooth performance for memory-intensive tasks.
- Allows for large batch sizes and complex model architectures in ML training.
- Supports multi-tasking and virtualized environments with ease.

4. Storage

○ 12 TB SSD

- High-speed SSD storage ensures quick data access, faster read/write speeds, and reduced latency.
- Ideal for large datasets and database management required for ML workloads.
- Provides ample storage for data-intensive applications.

5. Operating System Options

○ Windows Server 2022

- Enterprise-grade operating system with enhanced security, virtualization, and remote management capabilities.
 - Optimized for running Windows-based ML frameworks and applications.
 - **Proxmox Server**
 - Open-source virtualization platform for managing virtual machines and containers.
 - Suitable for deploying multiple isolated ML environments on a single physical server.
-

Key Features and Benefits

1. Performance

- The dual Intel Xeon processors and Tesla GPU provide exceptional computational power for training deep learning models and running simulations.
- High memory capacity supports large-scale ML models and in-memory databases.

2. Scalability

- Expandable RAM and storage allow the server to grow with increasing computational needs.
- Modular design supports additional GPUs and network cards.

3. Reliability

- Enterprise-grade hardware ensures uptime and consistent performance.
- Advanced cooling and power management for sustained operation under heavy workloads.

4. Flexibility

- Dual OS options (Windows Server 2022 or Proxmox) cater to diverse use cases.
- Compatible with a wide range of ML libraries and frameworks.

5. High-Speed Data Access

- The 12 TB SSD reduces data bottlenecks, ensuring seamless training and real-time inference.
-

Applications

1. Machine Learning and AI

- Training deep learning models, including image recognition, natural language processing, and predictive analytics.
- Real-time inference for AI-driven applications.

2. Data Analytics

- Big data processing and real-time analytics for business intelligence.

3. **Virtualization**

- Hosting multiple virtual machines for testing, development, or deployment.

4. **High-Performance Computing (HPC)**

- Running simulations, scientific computations, and complex algorithms.