

Micron® 9550 NVMe™ SSD + innovative storage software enable efficient AI model training offload



Artificial intelligence (AI) model sizes are growing rapidly and increasing in complexity.¹ One method for training extremely large models is to have as much high bandwidth memory (HBM) as possible on the GPU, along with as much system DRAM as possible. If a model doesn't fit in this HBM + DRAM, it can be parallelized over multiple GPU systems.

However, there is a heavy cost to parallelized training over multiple servers: lower GPU utilization and decreased efficiency due to data flow over network and system links. These can easily become bottlenecks.

What if splitting an AI training job over multiple GPU systems can be avoided by using NVMe storage as a third tier of “slow” memory? That's exactly what Big accelerator Memory (BaM) with GPU Initiated Direct Storage (GIDS) does. It replaces and streamlines the NVMe driver, handing the data and control paths to the GPU.

BaM and GIDS are research projects based on the following paper, with open-source code available on GitHub:

- [GPU-Initiated On-Demand High-Throughput Storage Access in the BaM System Architecture](#)
- [BaM open-source code \(GitHub\)](#)

The BaM software stack uses the low latency, extremely high throughput, large density, and high endurance of NVMe SSDs as a memory extension. BaM uses a custom storage driver that is optimized for the inherent parallelism of GPUs to access storage devices directly. No modifications to the SSD are necessary.

The GIDS data loader is built on the BaM subsystem to address memory capacity requirements for GPU-accelerated Graph Neural Network (GNN) training workload.

Using the Illinois Graph Benchmark (IGB) heterogeneous full dataset² in this testing showed the Micron 9550 SSD offers better training performance, lower system energy consumption, and strong scaling to help fully utilize the PCIe Gen5 x16 links of the NVIDIA® H100 Tensor Core GPU.

Key findings

As AI models have grown rapidly in size and complexity, training them requires more memory than may be available on GPUs and in training servers. BaM and GIDS use NVMe SSDs to extend memory.

Testing showed the Micron 9550 SSD enabled faster GNN training, demonstrated higher SSD performance, used less system energy,³ and provided strong scaling results.

33% **Faster training workload completion**

The Micron 9550 SSD showed workload execution time that was up to 33% faster than the competition. Training completes faster so systems can move on to the next training job faster, helping drive overall efficiency.

60% **Higher SSD performance**

Head-to-head testing demonstrated the Micron 9550 SSD offered up to 60% higher SSD performance during the GNN training workload compared to the competition. Higher SSD performance is at the heart of faster workload execution.

29% **Less energy used by training system**

A critical factor to consider in AI workloads is energy consumption. The training system that used the Micron 9550 SSD consumed up to 29% less energy to complete the same training workload.

Strong performance scaling

Scaling the number of Micron 9550 SSDs in the training system led to a substantial reduction in GNN workload completion times.

micron.com/9550

1. [Machine Learning Model Sizes and the Parameter Gap \(EPOCH AI\)](#).

2. See [Creating Large Real and Synthetic Graph Datasets for GNN Applications](#) for additional details.

3. Values are maximums observed during testing. Competitive PCIe Gen5 SSDs (Kioxia CM7-R and Samsung PM1743) chosen from the top 10 PCIe SSD suppliers shown in the Forward Insights analyst report “SSD Supplier Status Q1/24 May 2024.”

Micron 9550 SSD showed higher GNN training workload performance

Figure 1 represents the end-to-end GNN training workload completion time split into three steps: sample time, feature aggregation time, and training time. Total workload completion time, in seconds, is shown on the horizontal axis (smaller is better).

Feature aggregation: The process where distinctive characteristics (features) from the data are grouped together to help the AI system understand the data and improve its learning and predictions.⁴ This stage is heavily SSD dependent, and is represented in blue in Figure 1.⁵

Training and sampling: Training is the process where the machine learning model learns (tries to understand patterns in the data) from a dataset and makes predictions or decisions. It's an iterative process based on feedback and results.⁵ Sampling is the process of selecting a subset of data from a larger dataset. The selected data, or samples, are representative and unbiased, ensuring the model learns from an accurate subset of the dataset.⁶ Neither depend on SSD performance.

As seen in Figure 1, the Micron 9550 SSD was 21% faster than the Kioxia CM7-R and 33% faster than the Samsung PM1743. Since each test platform was identical (except for the SSD used), differences in workload completion time may be attributed to SSD performance.

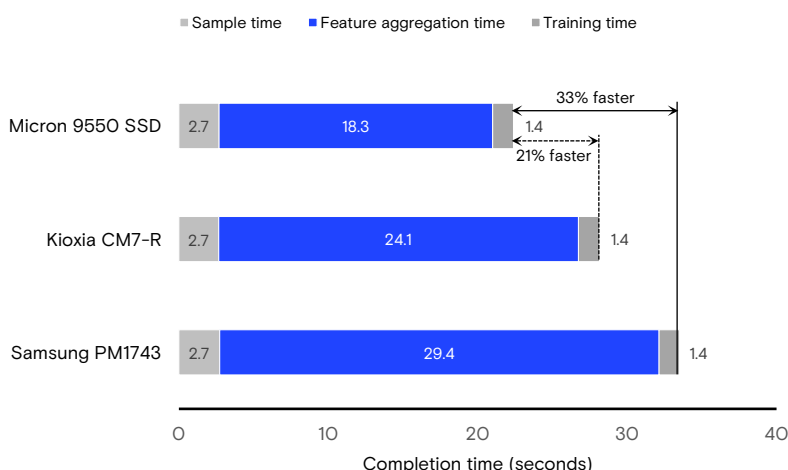


Figure 1: GNN workload completion time

A closer look at SSD performance in GNN workload testing

Figure 2 represents SSD performance during feature aggregation – the portion of the GNN workload testing that is highly SSD dependent.

IOPS are shown in solid rectangles while throughput (in GB/s) is shown in outlined rectangles.⁷

Figure 2 gives clear insight into why the Micron 9550 SSD GNN workload completion time is faster than either the competitors, as the Micron 9550 SSD performance was 31% higher than the Kioxia CM7-R and 60% higher than the Samsung PM1743.⁸

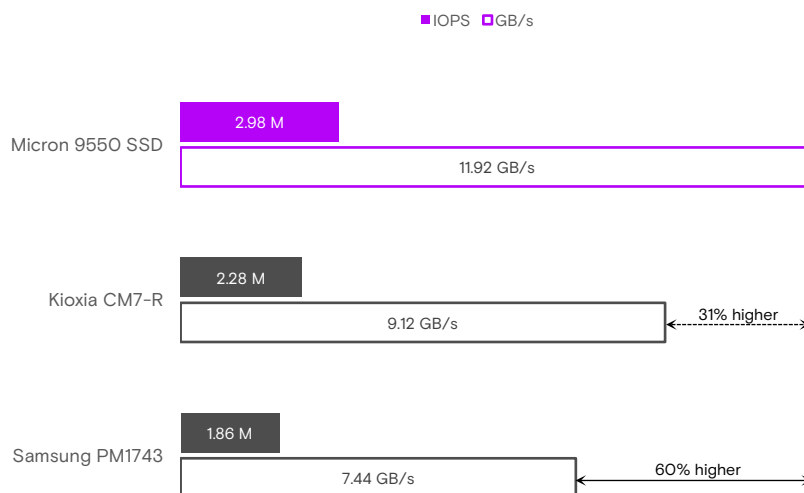


Figure 2: SSD performance during GNN workload testing, higher is better

4. See [Feature-Based Aggregation and Deep Reinforcement Learning](#) (Dimitri P. Bertsekas).

5. For brevity, platforms using a specific SSD are referred to only by the name of the SSD used.

6. See [What is Data Sampling and How is it Used in AI?](#) (Dataquest).

7. Note that purple fill and outline colors correspond to Micron 9550 SSD. This color correlation continues throughout other figures in this document.

8. Performance improvements are calculated as the percentage difference between the Micron 9550 SSD performance and competitor drives. Percentages for IOPS and throughput are similar.

Higher SSD performance lowers total system energy consumption

Beyond the GNN workload performance, a critical factor to consider is energy consumption, particularly for AI workloads. To illustrate the magnitude of this issue, a report by Schneider Electric projects that by 2028, AI tasks will consume approximately 4.3 gigawatts globally. This staggering figure is comparable to the total energy consumption of an entire country—Denmark, to be precise.⁹

Since the Micron 9550 SSD completed the GNN training more quickly, it also helped reduce system energy used relative to the Kioxia CM7-R and the Samsung PM1743. Reduced energy consumption and higher performance led to increased system energy efficiency, as reflected in the lower total system energy consumed using the Micron 9550 SSD (shown in Figure 3).

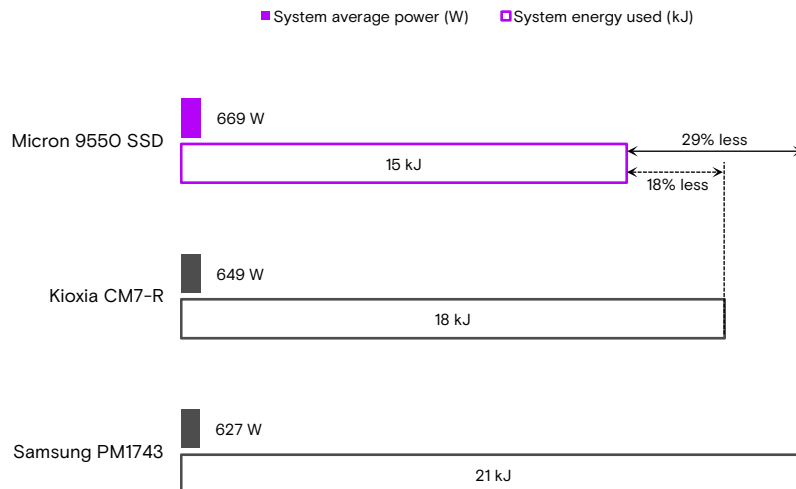


Figure 3: System energy used during GNN workload testing, lower is better

Figure 3 shows average system power (in watts) as a solid rectangle while system energy used (in kilojoules – note that 1 kilojoule = 1,000 watt-seconds) is reflected in an outlined rectangle, similar to Figure 2.

Using the Micron 9550, the test system consumed up to 29% less total energy training the GNN. Similar average power draw with faster workload completion translates into major energy savings.¹⁰

A closer look at SSD power and energy used in GNN workload testing

Figure 4 represents SSD average power (in watts) shown in solid rectangles and SSD energy used (in joules) shown in outlined rectangles.

The Micron 9550 SSD consumed significantly less energy during the GNN workload testing: 36% less energy than the Kioxia CM7-R and 43% less energy than the Samsung PM1743.¹¹

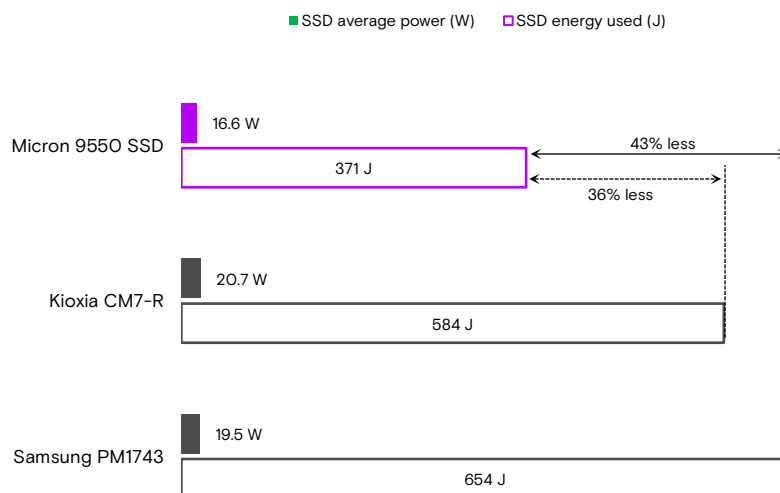


Figure 4: SSD average power and total energy used during GNN workload testing, lower is better

9. See [Schneider Electric predicts substantial energy consumption for AI workloads globally](#).

10. Differences are calculated as (Micron 9550 SSD value – competitor value) / competitor value, expressed as a percentage. Test systems were identical except for the SSD used, making system-level energy consumption differences attributable to the SSD.

11. SSD energy use improvements are calculated as ((competitive SSD energy used) - (Micron 9550 SSD energy used)) / (competitive SSD energy used), expressed as a percentage, or (653.8 - 371.0) / 653.8 = 43% less.

BaM and GIDS GNN training performance scales very well

BaM and GIDS enable the H100 GPU to utilize the Gen5 PCIe interface more fully due to the natural bandwidth scaling of multiple SSDs.

Figure 5 shows the GNN workload completion time (including sampling time, feature aggregation time, and training time) as more Micron 9550 SSDs were used. The number of Micron 9550 SSDs was scaled from one, to two, to four SSDs.

When the number of Micron 9550 SSDs was increased from one to two, the GNN workload completion time decreased from 22.4 seconds to 14.1 seconds, a 37% reduction. When the number of Micron 9550 SSDs was increased from one to four, the GNN workload completion time decreased 56%.¹²

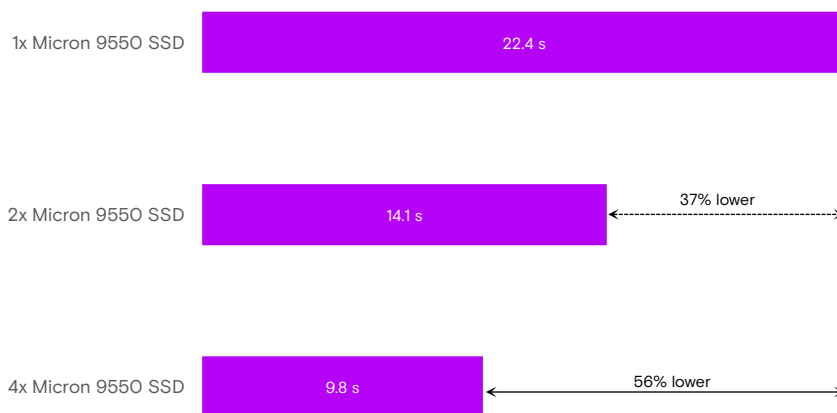


Figure 5: GNN workload completion time reduction by scaling Micron 9550 SSD count, shorter is better (s = seconds)

Scaling the number of Micron 9550 SSDs used can lead to better utilization of GPU resources, enabling improved performance of larger, more complex model training.

A close look at system energy when scaling Micron 9550 SSDs on GNN workloads

With the growing concern over the environmental impact of computing, reducing the energy consumption of training workloads is becoming increasingly important. Simply put, faster training times can contribute to lower energy use.¹³

Figure 6 represents system energy (in kilojoules) used during the GNN workload for each SSD configuration, shown in outlined rectangles. Although it may seem counter-intuitive, system energy use decreased when more Micron 9550 SSDs were added because training completed faster.

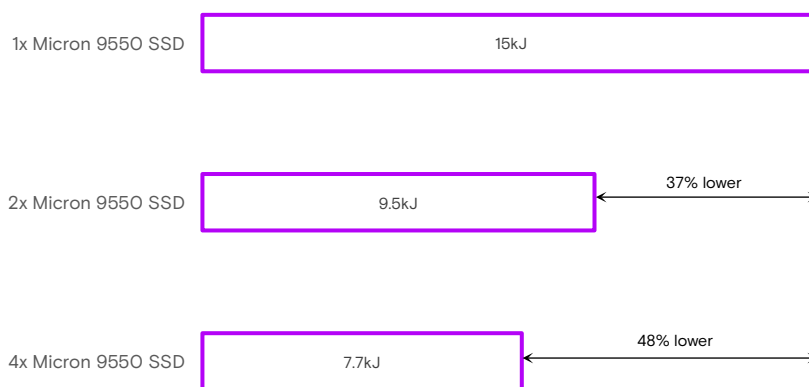


Figure 6: GNN workload energy use when scaling Micron 9550 SSD count, smaller is better

When scaling from one Micron 9550 SSD to two, system energy use decreased 37%. When scaling from one to four, system energy use decreased 48%.¹⁴

12. Percentage decrease from one SSD to two SSDs calculated as $((1x \text{ Micron } 9500 \text{ SSD GNN workload completion time}) - (2x \text{ Micron } 9500 \text{ SSD GNN workload completion time})) / ((1x \text{ Micron } 9500 \text{ SSD GNN workload completion time}) - (2x \text{ Micron } 9500 \text{ SSD GNN workload completion time}))$, or $(22.4 - 14.1) / 22.4 = 37\%$ less. The percentage decrease from one to four is calculated similarly.

13. See [CoGNN: Efficient Scheduling for Concurrent GNN Training on GPUs](#) for additional details on GNN efficiency.

14. Percentage use decrease calculated using method described in footnote 12.

Conclusion

AI model sizes are growing rapidly. Innovative software solutions like BaM and GIDS are becoming essential elements to satisfy use cases where GPU-enabled hardware is at a premium.

Comparing the Micron 9550 SSD to competitive SSDs showed that a training system using the Micron 9550 SSD offered up to 33% faster GNN training workload completion time, 60% higher SSD performance, and decreased system energy consumption by up 29%.

Additionally, scaling the number of Micron 9550 SSDs from one to two showed that the GNN training workload completion time decreased by 37% and system energy consumption decreased by the same amount. Scaling the number of Micron 9550 SSDs from one to four decreased GNN training workload completion time by 56% and lowered the system energy used by 48%.

How we tested

Tables 1 and 2 outline the system and software configurations used.

Hardware configuration	
GPU	NVIDIA H100 NVL GPU (94GB, driver 535.161.08)
Server	Supermicro® SuperServer SYS-521GE-TNRT
CPU s	2X Intel® Xeon® Platinum 8568Y+ processors (48 cores per CPU, 192 total vCores with Hyperthreading)
Memory	Micron 96GB, DDR5 5600MT/s (x16 = 1.5TiB)
Micron SSDs	Micron 9550 NVMe SSD, 7.68TB
Competitors' SSDs	Kioxia CM7-R, Samsung PM1743 (7.68TB advertised capacity for each SSD)

Table 1: Test platform configuration

Software configuration	
OS	Ubuntu 20.04 LTS (5.4.0-182-generic kernel)
Kernel	5.4.0-182-generic
CUDA	12.4
NVIDIA driver	535.161.08

Table 2: Software configuration

micron.com/9550

© 2024 Micron Technology, Inc. All rights reserved. All information herein is provided on as “AS IS” basis without warranties of any kind, including any implied warranties, warranties of merchantability or warranties of fitness for a particular purpose. Micron, the Micron logo, and all other Micron trademarks are the property of Micron Technology, Inc. All other trademarks are the property of their respective owners. Products are warranted only to meet Micron’s production data sheet specifications. Products, programs and specifications are subject to change without notice. Rev. A 07/2024 CCM004-676576390-11759