

NeuralMesh™ Accelerates Autonomous Vehicle Training, Simulation, and Fleet Ingest

WEKA's NeuralMesh™ eliminates storage bottlenecks across sensor ingest, model training, data curation, and simulation to keep GPU clusters fed, accelerate AV development cycles, and support safety compliance for autonomous vehicle organizations.

Key Benefits of NeuralMesh

Reproduce Any Training Run on Demand

- Version every dataset from ingest to production
- Trace lineage to sensor frame
- Support regulatory defense

Checkpoint Training Runs Without Stalls or Re-Runs

- Absorb checkpoint write storage constantly
- Complete training cleanly
- Protect expensive GPU cycles

Accelerate Simulation and Synthetic Data

- Cover gaps that real on-road driving misses
- Replay synthetic sensor data
- Generate rare edge cases fast

Maximize GPU Saturation Across AV Workloads

- Feed complex training clusters nonstop
- Sustain 90%+ GPU utilization
- Eliminate idle GPU cycles

Unify Storage across Edge, On-Premises, and Cloud

- Operate on all data in a single namespace
- Share sensor & training data
- Eliminate cross-tier copies

Handle Billions of Small Sensor Files

- Promote archive data to hot tier on demand
- Scale past metadata limits caused by lots of small files

Integrate Directly with NVIDIA Platforms

- Validate training runs on DGX SuperPOD
- Use GPUDirect Storage
- Leverage NVMe-to-GPU I/O

Absorb Fleet Ingest Without Backlogs

- Clear pre-processing backlogs automatically
- Ingest petabyte fleet data
- Route data to training lake

>200,000

GPUs in a single NeuralMesh cluster

10x

More bandwidth compared to Legacy NAS

14 Days

Reduced to just **4 hours** for a single epoch

>95%

GPU utilization, up from ~30%

600+ PB

Of data supported at 16TB/s all day, every day

Autonomous Vehicle Use Cases

Sensor Data Capture and Ingest

NeuralMesh eliminates pre-processing backlogs by ingesting PB-sized fleet data from LiDAR, radar, and camera sensors directly into the training data lakes

AV Model Training and Validation

NeuralMesh delivers concurrent, high-bandwidth access across distributed workloads with frequent checkpointing, eliminating I/O waits and stalls

Data Curation and Active Learning

NeuralMesh resolves LOSF challenges and enables instant promotion of archive data back into the active training set for faster iteration

Simulation and Synthetic Data Generation

NeuralMesh feeds simulation workloads at the throughput required to produce, store, and replay synthetic sensor data for rare-scenario edge cases

“WEKA’s storage scalability and ability to grow the infrastructure without losing performance was the key factor.”

- Oren Ben Ibghei, IT manager, Innoviz

