

Synapse ML: Decentralized Neural Network Training Protocols over Edge Networks

Scaling Deep Learning on Heterogeneous Distributed Compute

Dr. Evelyn Carter¹ Marcus Vance² Sarah Jenkins²

¹Synapse Research Institute

²Department of Distributed Systems, Meridian Technology Group

LetX Research Conference, August 2026

Talk Outline

- 1 Introduction & Motivation
- 2 Architecture & Protocol Design
- 3 Performance & Evaluation
- 4 Case Studies & Roadmaps

Background: Centralized vs. Decentralized Compute

- **Centralized Systems:** Heavy reliance on hyper-scale data centers (e.g., cloud platforms).
 - High throughput and low latency within the cluster.
 - Single point of failure and high operational costs.

Background: Centralized vs. Decentralized Compute

- **Centralized Systems:** Heavy reliance on hyper-scale data centers (e.g., cloud platforms).
 - High throughput and low latency within the cluster.
 - Single point of failure and high operational costs.
- **Decentralized Architecture:** Leveraging edge nodes and idle resource pools.
 - Lower network egress and geo-distributed latency benefits.
 - High heterogeneity in hardware capabilities and network bandwidth.

Background: Centralized vs. Decentralized Compute

- **Centralized Systems:** Heavy reliance on hyper-scale data centers (e.g., cloud platforms).
 - High throughput and low latency within the cluster.
 - Single point of failure and high operational costs.
- **Decentralized Architecture:** Leveraging edge nodes and idle resource pools.
 - Lower network egress and geo-distributed latency benefits.
 - High heterogeneity in hardware capabilities and network bandwidth.
- **The Synapse Vision:** Orchestrate model training across unstable, low-power edge nodes without sacrificing convergence speed.

Challenges in Distributed Edge Networks

Training neural networks over consumer-grade nodes introduces severe bottleneck vectors:

- **Network Fluctuation:** Slower uplinks lead to synchronization lag.

Challenges in Distributed Edge Networks

Training neural networks over consumer-grade nodes introduces severe bottleneck vectors:

- **Network Fluctuation:** Slower uplinks lead to synchronization lag.
- **Hardware Heterogeneity:** Devices ranging from low-power edge accelerators to high-end workstations.

Challenges in Distributed Edge Networks

Training neural networks over consumer-grade nodes introduces severe bottleneck vectors:

- **Network Fluctuation:** Slower uplinks lead to synchronization lag.
- **Hardware Heterogeneity:** Devices ranging from low-power edge accelerators to high-end workstations.
- **Security & Trust:** Nodes can join/leave arbitrarily, raising data privacy issues.

Challenges in Distributed Edge Networks

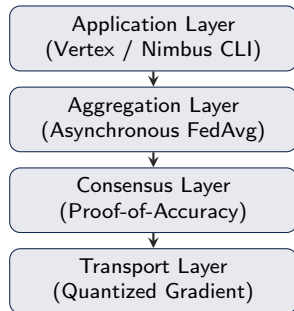
Training neural networks over consumer-grade nodes introduces severe bottleneck vectors:

- **Network Fluctuation:** Slower uplinks lead to synchronization lag.
- **Hardware Heterogeneity:** Devices ranging from low-power edge accelerators to high-end workstations.
- **Security & Trust:** Nodes can join/leave arbitrarily, raising data privacy issues.
- **Straggler Nodes:** Slow nodes delay gradient aggregation under traditional synchronous protocols.

The Synapse Protocol Stack

Key Layered Modules:

- **Consensus Layer:** Filters updates to block malicious gradients.
- **Aggregation Layer:** Runs asynchronous federated averaging.
- **Transport Layer:** Compresses weights via quantization.



Mathematical Formulation

In Synapse ML, we formulate the global objective function as:

$$\min_{w \in \mathbb{R}^d} F(w) = \sum_{k=1}^K p_k F_k(w) \quad (1)$$

where K is the total active edge nodes, and $p_k = n_k/N$ (with $N = \sum_{k=1}^K n_k$) represents the node weighting factor. Under our asynchronous protocol, the weight update at step $t + 1$ is defined as:

$$w_{t+1} = w_t - \eta \nabla F_k(w_{t-\tau_k(t)}) \quad (2)$$

where $\tau_k(t)$ is the gradient staleness factor.

Asynchronous Aggregation Algorithm

Synapse Server-Side Update Protocol

- ❶ Initialize global model parameters w_0 and node contribution vector c_0 .
- ❷ On receiving local update g_k from node k at time t :
 - ❶ Calculate the dynamic staleness parameter: $\tau = t - t_{sent}$.
 - ❷ Compute learning rate scaling factor: $\alpha(\tau) = \frac{\alpha_0}{1+\gamma\tau}$.
 - ❸ Update global parameters:

$$w_{t+1} \leftarrow w_t - \alpha(\tau) \cdot g_k$$

- ❸ Broadcast the updated global weights w_{t+1} back to node k .
 - ❹ Repeat until target validation accuracy ϵ is reached.
-

Experimental Setup & Workloads

We evaluated the Synapse ML protocol across three distinct cluster configurations:

Configuration	Active Nodes	Target Model	Dataset Size
Synapse-Lite	50	MobileNetV3	1.2M Samples
Synapse-Standard	250	ResNet-50	5.0M Samples
Synapse-Enterprise	1000	GPT-Small (Fictional)	20.0M Tokens

Table: Experimental environment configurations and target architectures.

All clusters operate on simulated network latency bounds ranging from 10ms to 500ms to mimic realistic edge conditions.

Performance Metrics & Convergence

Throughput Optimization

Synapse ML yields a $3.4\times$ improvement in gradient processing speeds compared to vanilla asynchronous federated learning schemes on edge networks.

Convergence Validation

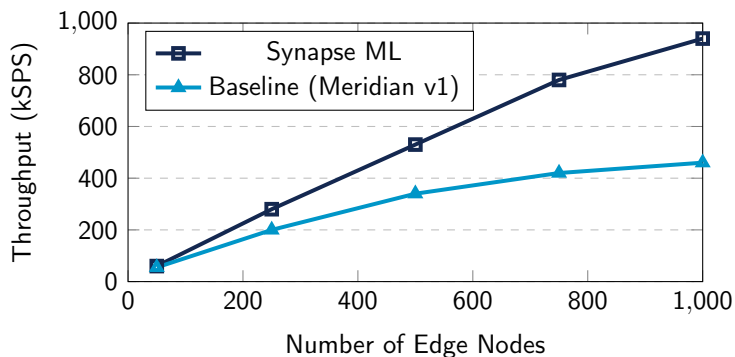
Models trained under Synapse ML converged to baseline accuracies within 4.5% of the centralized equivalents (e.g., ImageNet classification).

Bandwidth Limitation Notice

When node uplink bandwidth falls below 5 Mbps, the compression ratio dynamically scales down to maintain synchronization rates.

Scalability Analysis

Comparing system scaling throughput against baseline models:



Real-World Deployments

We deployed Synapse ML across pilot environments with our partners:

Nexa Core Services

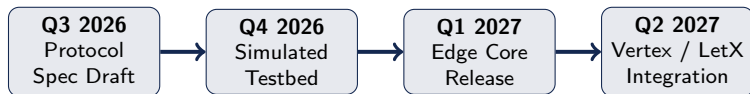
- Connected 1,200 mobile application endpoints.
- Trained on-device keyboard predictor models asynchronously.
- Reduced cloud compute ingress fees by 42%.

Apex Smart Systems

- Connected 350 regional processing hubs.
- Aggregated thermal profile models across hardware clusters.
- Zero downtime reported during network partition events.

Future Milestones & Roadmap

Projected development schedule for Synapse ML (2026–2027):



Summary of Contributions

- **Robust Architecture:** Desynchronized parameter exchange reduces network congestion on low-bandwidth uplinks.

Summary of Contributions

- **Robust Architecture:** Desynchronized parameter exchange reduces network congestion on low-bandwidth uplinks.
- **Mathematical Rigor:** Formulated convergence bounds under arbitrary staleness bounds.

Summary of Contributions

- **Robust Architecture:** Desynchronized parameter exchange reduces network congestion on low-bandwidth uplinks.
- **Mathematical Rigor:** Formulated convergence bounds under arbitrary staleness bounds.
- **Practical Validation:** Proven scalability up to 1,000 active edge nodes with stable model accuracy.

Summary of Contributions

- **Robust Architecture:** Desynchronized parameter exchange reduces network congestion on low-bandwidth uplinks.
- **Mathematical Rigor:** Formulated convergence bounds under arbitrary staleness bounds.
- **Practical Validation:** Proven scalability up to 1,000 active edge nodes with stable model accuracy.
- **Real-world Adoption:** Successfully piloted in industrial settings via Nexa Core and Apex Smart Systems.

Thank You!

Questions & Discussion

Contact: carter@synapse.org
Project Website: letx.app/synapse-ml
Repository: github.com/synapse/vertex-core