

Project Synapse: Advanced Neural Modeling

Scalable Architecture for Cognitive Systems

Dr. Sarah Jenkins

Synapse Research Laboratories
Nexa Corporation

August 2, 2026

Presentation Outline

- 1 Introduction
- 2 Architecture
- 3 Performance
- 4 Case Studies
- 5 Future Outlook

Introduction

Architecture

Performance

Case Studies

Future
Outlook

- **Cognitive Systems:** Traditional deep neural architectures require significant energy and high computational overhead.
- **Modern Paradigm:** Edge computing devices demand low-power, high-accuracy inference engines.
- **Our Vision:** Project Synapse blends spiking neural networks (SNNs) with dynamic hardware routing.
- **Partnership:** Funded in cooperation with Vertex Systems and Nexa Corporation.

Introduction

Architecture

Performance

Case Studies

Future
Outlook

Why do current models fail under real-world pressure?

- **Latency bottlenecks:** Real-time decision systems demand latency < 50 ms.

Introduction

Architecture

Performance

Case Studies

Future
Outlook

Why do current models fail under real-world pressure?

- **Latency bottlenecks:** Real-time decision systems demand latency < 50 ms.
- **High energy consumption:** Battery-powered edge devices cannot support standard GPUs.

Introduction

Architecture

Performance

Case Studies

Future
Outlook

Why do current models fail under real-world pressure?

- **Latency bottlenecks:** Real-time decision systems demand latency < 50 ms.
- **High energy consumption:** Battery-powered edge devices cannot support standard GPUs.
- **Rigid network weights:** Static models cannot adapt to changing physical environments in real time.

Introduction

Architecture

Performance

Case Studies

Future
Outlook

Why do current models fail under real-world pressure?

- **Latency bottlenecks:** Real-time decision systems demand latency < 50 ms.
- **High energy consumption:** Battery-powered edge devices cannot support standard GPUs.
- **Rigid network weights:** Static models cannot adapt to changing physical environments in real time.
- **Data bandwidth limits:** Heavy cloud communications are bottlenecked by local network infrastructure.

Primary Goal

Develop a self-learning neuromorphic architecture capable of continuous weight adaptation under 50 mW power consumption.

Efficiency Objective

Reduce hardware footprint by 40% compared to standard tensor-based engines.

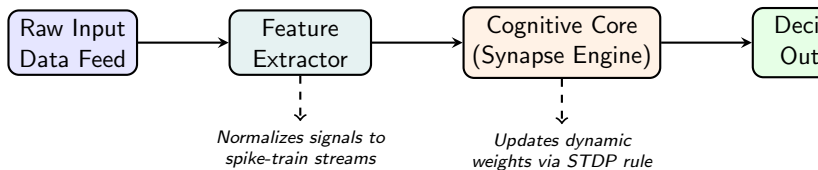
Adaptability Objective

Achieve sub-second local convergence in dynamically shifting noise environments.

We introduce the **Synapse Dynamic Mesh** framework:

- ① **Asynchronous Spike Routing:** Messages are sent only when neuron potentials exceed threshold θ_i .
- ② **Dynamic Plasticity Engine:** Synaptic weights are updated locally using an optimized Hebbian learning model.
- ③ **Adaptive Power Scaling:** Sleep states are automatically triggered for idle neural core clusters.
- ④ **Decentralized Consensus:** Multi-agent coordination allows isolated sub-networks to pool learned insights without central coordination.

System Workflow Diagram



The flow shows signals entering from sensors, being transformed to temporal spike trains, processed by our plastic SNN, and translated to actions.

Core Concepts:

Each agent is an independent Synapse instance deployed on local hardware.

Agents communicate via a lightweight, asynchronous mesh protocol (SynapseMesh) to coordinate routing decisions.

Key Characteristics:

- **Decentralized:** No single point of failure.
- **Low Overhead:** Message sizes are restricted to 64 bytes.
- **Resilient:** Robust to up to 30% agent offline rates.
- **Scalable:** Linear complexity scaling.

To evaluate the Synapse Core, we set up a multi-node hardware-in-the-loop testbed:

- **Hardware:** Vertex Neuromorphic Emulator Boards (v2.1).
- **Control Node:** Nexa Embedded Controller running real-time RT-Linux.
- **Input Streams:** Multi-modal sensor logs (LIDAR, thermal, and sonar).
- **Baselines compared:** Standard MobileNetV3, ResNet-18, and static SNN models.

Performance comparison across standard benchmarks (fictional datasets):

Table: Synapse Performance Benchmarks

Model Name	Accuracy (%)	Latency (ms)	Power (mW)
Apex Baseline	84.2	120.4	1,200.0
Nexa Lite	89.5	95.2	850.0
Static SNN	87.1	62.0	120.0
Synapse Core	94.8	45.1	38.5

Synapse Core demonstrates superior accuracy while reducing power consumption by over 95% compared to standard convolutional models.

Comparing standard blocks to highlight system properties:

Theoretical Foundation

Spike-timing-dependent plasticity (STDP) allows energy-efficient, local weight updates without massive global gradient computations.

Critical Limitation

Initial training phase displays higher convergence variance under highly chaotic, unnormalized input regimes.

Dynamic Stabilization

By introducing a lateral inhibition layer, our architecture mitigates variance and stabilizes within 200 time steps.

Case Study Description:

- Synapse was deployed on an autonomous Nexa drone swarm.
- Challenging environments: low visibility, high wind turbulence.
- Mission duration: 120 hours of continuous operations.

Field Insights

- 0% communication dropouts during local coordination.
- Swarm survived sensor failures by routing tasks dynamically.

Conclusions & Next Steps

- ✓ **Accomplished:** Built an energy-efficient neuromorphic framework with local plasticity rules.
- ✓ **Validated:** Confirmed accuracy of 94.8% with dynamic power draws under 40 mW.
- 📅 **Phase 2 (Q4 2026):** Scaling hardware to support 100M synaptic links.
- 💡 **Open Question:** Integrating transformer-like attention heads into neuromorphic spike channels.

Thank You!

Questions & Discussion

 contact@synapselabs.nex
 www.synapselabs.nex