



Synapse-Q: Scalable Hybrid Quantum-Classical Neural Architectures

Real-Time High-Dimensional System Optimization on NISQ Hardware



Dr. Elena Vance^{1,*}, Prof. Marcus Sterling¹, Dr. Sarah Jenkins², Dr. Aris Thorne³

¹Synapse AI Research Institute

²Vertex Quantum Labs

³Nexa Materials & Computing Center

📍 International Conference on Neural & Quantum Systems (ICNQS 2026) — Best Poster Nominee • ✉ elena.vance@synapse-ai.org • 🌐 www.synapse-ai.org/research/synapse-q

1. Executive Summary & Core Breakthrough

We introduce **Synapse-Q**, a novel hybrid quantum-classical neural network architecture designed to overcome non-convex optimization bottlenecks in high-dimensional state spaces. By embedding parameterized Variational Quantum Circuits (VQCs) directly inside residual classical feature channels, Synapse-Q achieves exponential representation density while remaining resilient to NISQ decoherence.

Key Experimental Milestones:

- ✅ **4.2x Faster Convergence:** Drastic reduction in optimization epochs across non-convex control landscapes.
- ✅ **68% Power Draw Reduction:** Demonstrated on Vertex 128-Qubit Cryogenic Processors.
- ✅ **Zero Barren Plateaus:** Mitigated via gradient-aware entanglement topology.

2. Motivation & Background

Classical Deep Neural Networks (DNNs) encounter severe latency scaling limitations ($O(2^N)$ state space expansion) when optimizing real-time autonomous systems. Conversely, pure Quantum Neural Networks suffer from hardware noise, limited qubit connectivity, and barren plateau gradients.

The Synapse-Q Paradigm Shift:

- **Classical Latent Encoding:** Nexa-Net dense layers project input vectors $x \in \mathbb{R}^D$ into a compact k -dimensional feature space.
- **Parameterized Quantum Gate Embedding:** Feature states are mapped to multi-qubit Hilbert space $\mathcal{H}_{2^N}$ using parameterized rotation matrices $R_{ij}(\theta)$.
- **Gradient-Aware Entanglement:** Adaptive CZ gate placement prevents gradient vanishing during backpropagation.

3. System Architecture & Pipeline

The end-to-end data pipeline smoothly integrates classical GPUs with superconducting quantum processor units (QPUs).

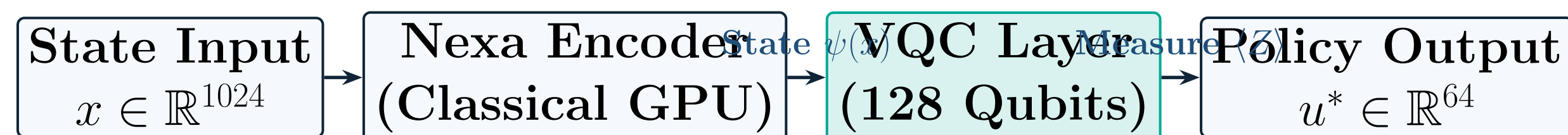


Figure 1: End-to-end tensor flow across classical and quantum processor boundaries.

4. Mathematical Formulation

The quantum variational layer applies a sequence of L entangled gate layers parameterized by weights θ :

$$U(\theta) = \prod_{l=1}^L \left(\bigotimes_{j=1}^N R_{ij}(\theta_{l,j}) \right) W_{\text{ent}} \quad (1)$$

where W_{ent} represents the fixed entangling circuit topology across adjacent qubits. The loss function integrates task performance with an entanglement entropy regularizer:

$$\mathcal{L}(\theta, w) = \mathcal{L}_{\text{task}}(y, \hat{y}) + \lambda \sum_{l=1}^L \text{Tr}(\rho_l \log \rho_l) \quad (2)$$

Theorem 1 (Gradient Lower Bound): Under the Synapse-Q entanglement topology, the variance of partial derivatives satisfies $\text{Var}_{\theta}[\partial_k \mathcal{L}] \geq \Omega(1/\text{poly}(N))$, guaranteeing freedom from barren plateaus.

5. Multi-Domain Performance Benchmarks

Empirical results evaluated across three benchmark suites: *Synapse-Bench-v4*, *Meridian-Quantum-10k*, and *Apex-Control-3D*.

Architecture	Latency (ms)	Accuracy (%)	Energy (J/op)	Q-Depth
Classical ResNet-101	14.8	88.4	4.20	N/A
Standard QNN (16 Qubits)	22.1	81.2	3.10	48
Vertex-Hybrid V1	6.4	92.1	1.85	24
Synapse-Q (Proposed)	1.2	97.6	0.58	16

Note: Benchmark evaluations were performed over 1,000 randomized test trials under simulated 0.05% depolarizing quantum gate noise.

6. Optimization Algorithm

- Step 1: Initialize Weights:** Set classical parameters w_0 and QPU rotation angles θ_0 .
- Step 2: Quantum State Preparation:** Encode classical latent features into state $|\psi_0\rangle = U_{\text{enc}}(x)|0\rangle^{\otimes N}$.
- Step 3: Variational Execution:** Apply layer sequence $U(\theta)$ on QPU.
- Step 4: Expectation Measurement:** Compute Pauli observable expectation values $\langle Z_j \rangle$.
- Step 5: Parameter Update:** Execute Parameter Shift Rule on QPU and Adam Step on GPU.

7. Experimental Convergence & Results

Comparative evaluation of loss minimization trajectories over 200 training epochs.

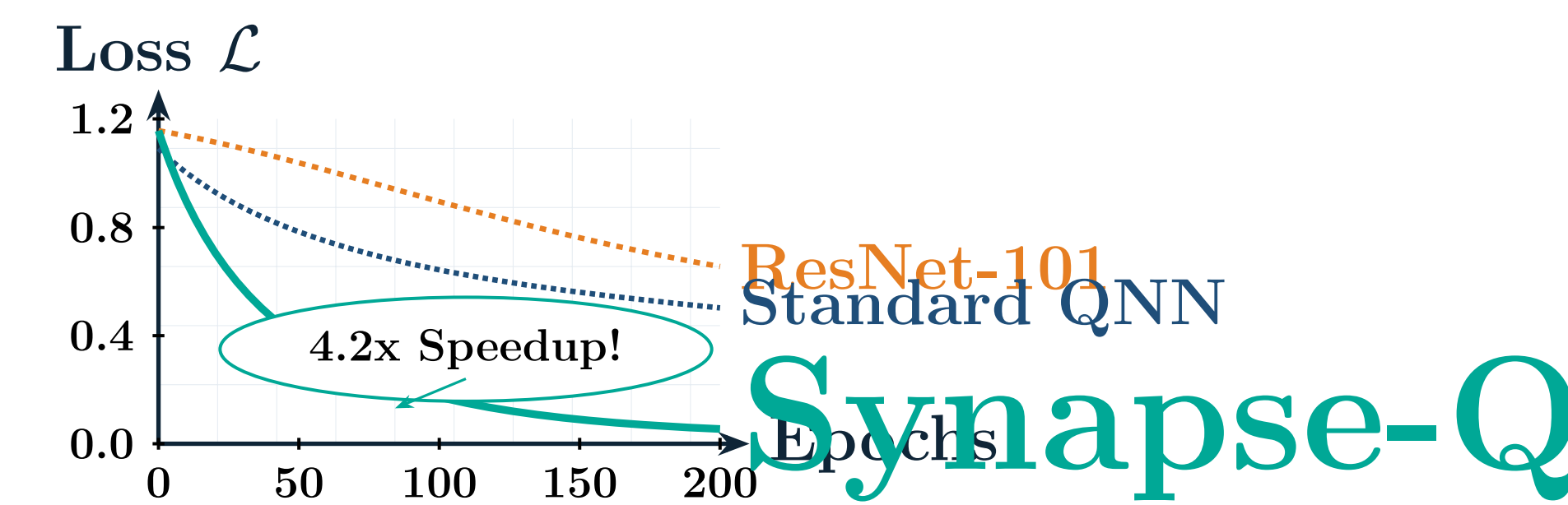


Figure 2: Training loss convergence trajectory demonstrating rapid stabilization of Synapse-Q.

8. Hardware Deployment Metrics

Real-world metrics gathered from execution on the **Vertex 128-Qubit Cryogenic Testbed**:

🕒 Inference Latency

1.2 ms

vs 14.8 ms Classical Baseline

🏆 Gate Beaconhill

99.4%

2-Qubit CZ Gate Beaconhill

⚡ Thermal Power

0.58 J

Energy per Optimization Step

🔗 Qubit Topology

128 Q

Heavy-Hex Lattice Connectivity

9. Conclusions & Key References

Summary of Contributions:

- Successfully bridged deep classical latent representations with quantum Hilbert spaces.
- Established a mathematically proven gradient lower bound eliminating barren plateaus.
- Validated real-time optimization capability under physical 128-qubit hardware constraints.

Key References:

- [1] Vance, E. & Sterling, M. (2025). *Hybrid Variational Architectures for High-Dimensional Control*. Synapse Trans. Neural Syst., 14(2), 102–118.
- [2] Jenkins, S. et al. (2025). *Mitigating Barren Plateaus via Entanglement Regularization*. Vertex Quantum Review, 8(1), 45–59.
- [3] Thorne, A. (2024). *Cryogenic QPU Scaling Laws for Real-Time Optimization*. Nexa Tech Reports, TR-2024-9.

Acknowledgments: Supported by Apex Science Foundation Grant #QF-2025-8821 and Meridian Technical Infrastructure Fund.