

ACCELERATING LARGE LANGUAGE MODEL TRAINING USING HYBRID QUANTUM-CLASSICAL VARIATIONAL AUTOENCODERS WITH INTEGRATED PHOTONIC QRAM

Pawan Kumar Singh

Dr. Pillai Global Academy, New Panvel, Sector-7, Khanda Colony, New Panvel, Navi Mumbai- 410206,
pawan.k.singh2009@gmail.com

Received: 12/06/2026

Revised: 19/07/2026

Accepted: 14/08/2026

ABSTRACT:

Training large language models (LLMs) has become one of the most computationally expensive undertakings in modern science and industry. State-of-the-art models require weeks of runtime on tens of thousands of accelerators, and the associated energy footprint is now large enough to attract serious environmental scrutiny. This paper explores whether hybrid quantum-classical variational autoencoders (QC-VAEs), paired with integrated photonic quantum random access memory (QRAM), can meaningfully accelerate parts of the LLM training pipeline. The premise is that certain latent-space operations, particularly compression, sampling, and gradient estimation over structured distributions, may benefit from quantum superposition and interference. We designed a hybrid architecture that offloads variational encoder and decoder branches to a parameterized quantum circuit, retrieves training minibatch representations through photonic QRAM with logarithmic addressing depth, and interleaves classical transformer training with periodic quantum-assisted latent refinement steps. Experiments were conducted through a co-simulation framework combining a classical GPU cluster model, a noisy intermediate-scale quantum (NISQ) device simulator, and a photonic QRAM performance model, evaluated on a 1.3-billion-parameter transformer trained on a curated 40-billion-token corpus. Results show that the hybrid pipeline reduced wall-clock training time by 27.4 percent and total energy consumption by 31.6 percent compared to a matched classical baseline, while final perplexity was within 0.8 percent of the classical result. Ablations revealed that photonic QRAM addressing latency and quantum circuit depth were the dominant sensitivity parameters. We discuss trade-offs including hardware maturity, noise tolerance, and integration cost, and outline a realistic path from co-simulation to early hardware deployment.

Keywords: *Large Language Models, Quantum Machine Learning, Variational Autoencoder, Photonic Qram, Hybrid Quantum-Classical Computing, Training Acceleration*

INTRODUCTION

The scale of modern large language model training has grown at a rate that no other computational workload has matched. Frontier models now involve trillions of parameters, are trained on tens of trillions of tokens, and consume computational budgets measured in exaflop-days. The direct cost of a single frontier training run is now in the tens to hundreds of millions of dollars, and the energy footprint of these runs has become a serious concern for both operators and the broader public [1]. Every meaningful reduction in training time or energy translates into real value, and the field has become intensely focused on any technique that can move the efficiency frontier. Most of the efficiency effort so far has focused on classical improvements. Mixed-precision arithmetic, sparse attention mechanisms, more efficient parallelism strategies, better data curation, and improved learning rate schedules have all contributed to steady gains [2]. Hardware improvements from GPU generation to GPU generation have added roughly a factor of two per cycle, and specialized AI accelerators have pushed further in specific regimes. Yet the scaling laws that govern how model quality improves with compute have not become friendlier, and each new generation of frontier models continues to require substantially more compute than its predecessor.

Quantum computing has long been proposed as a potential source of algorithmic speedup for machine learning tasks, though the picture has been genuinely mixed. Some proposed quantum machine learning algorithms have turned out to have efficient classical analogs once carefully analyzed, while others depend on assumptions about data loading that are not currently realistic [3]. Yet certain families of quantum algorithms, particularly variational

ones that combine parameterized quantum circuits with classical optimization loops, have shown genuine promise for structured optimization and sampling tasks that appear inside modern deep learning pipelines [4]. Variational autoencoders in particular have a natural quantum analog because their latent-space distributions and sampling operations map well to the state preparation and measurement primitives of quantum circuits.

A separate but related development is the maturation of integrated photonic quantum random access memory (QRAM). Traditional QRAM proposals suffered from decoherence problems that made large-scale implementation impractical. Recent work on photonic implementations, in which addressing paths are routed through integrated optical circuits rather than through long chains of quantum gates, has revived interest in QRAM as a practical component of near-term quantum systems [5]. Because photonic QRAM can address exponentially many locations with logarithmically many switching operations, it becomes plausible to load structured data into quantum states quickly enough for machine learning applications.

This paper asks whether combining these two developments, a hybrid quantum-classical variational autoencoder and integrated photonic QRAM, can meaningfully accelerate large language model training. The core intuition is that certain operations inside modern LLM training pipelines, particularly latent-space compression and sampling for regularization or distillation, are exactly the kind of tasks where quantum circuits could plausibly outperform classical alternatives at moderate qubit counts. Our research questions are: how much wall-clock and energy improvement can a realistic hybrid quantum-classical pipeline provide for LLM training, how does the improvement decompose across the architectural components involved, and what practical barriers stand between the current state of the art and a deployable system.

LITERATURE REVIEW

The intersection of quantum computing and machine learning has been actively studied for more than a decade. Early theoretical work outlined potential exponential speedups for specific linear algebra primitives such as inversion of well-conditioned matrices, but subsequent analysis showed that many of these speedups either depended on classical preprocessing assumptions that eliminated the advantage or admitted efficient classical alternatives [6]. This tempered but did not extinguish interest in quantum machine learning, and attention shifted toward more modest but achievable goals.

Variational quantum algorithms have become the dominant paradigm for near-term quantum machine learning. These algorithms use parameterized quantum circuits, sometimes called quantum neural networks, whose parameters are optimized classically to minimize a cost function evaluated on measurement outcomes [7]. Variational quantum eigensolvers for chemistry, variational quantum classifiers for structured data, and quantum generative adversarial networks have all been demonstrated in simulation and, at small scale, on hardware. Studies have shown that these methods can offer advantages in specific structured problems, though clear separation from classical baselines at scale remains elusive [8].

Quantum variational autoencoders have received particular attention because their structure aligns naturally with quantum information processing. The encoder maps input data to a latent quantum state, an entangled representation is manipulated by a parameterized quantum circuit, and a decoder measurement yields a reconstruction [9]. Applied to structured data such as molecular representations or graph embeddings, quantum VAEs have shown competitive results with fewer parameters than classical baselines, though at costs that are still difficult to justify at production scale.

Photonic quantum computing has emerged as one of the leading platforms for scalable quantum information processing. Compared to superconducting qubits, photonic systems offer room-temperature operation, natural connectivity through waveguides, and compatibility with existing semiconductor manufacturing infrastructure [10]. Photonic QRAM specifically has attracted interest because photonic routing is naturally suited to the tree-structured addressing patterns required by RAM-like access to quantum-encoded data. Recent proposals for integrated photonic QRAM using thin-film lithium niobate and silicon nitride waveguides have shown promising theoretical performance at scales relevant to machine learning applications [11].

The intersection of quantum acceleration and deep learning training has been explored in a few directions. Some work has focused on quantum optimizers that could replace or augment classical optimizers such as Adam. Others have proposed quantum-assisted attention mechanisms in which portions of the attention computation are

offloaded to quantum circuits [12]. Studies of quantum-assisted gradient estimation have shown modest theoretical advantages for structured loss landscapes, though experimental validation at meaningful scale remains limited [13].

Hybrid architectures that combine classical GPUs with quantum coprocessors are widely viewed as the most practical near-term paradigm. In this model, the classical hardware handles the vast bulk of the computation, and the quantum coprocessor handles specific subroutines where it offers advantages. Studies of hybrid architectures for chemistry, optimization, and small-scale machine learning have shown that the interconnect and control overhead can be managed if the quantum subroutine execution time is long enough to amortize the setup cost [14]. Application of this paradigm specifically to LLM training has been discussed conceptually but not extensively evaluated, which motivates the present study [15].

METHODOLOGY

Our hybrid architecture combines a classical transformer training pipeline with a quantum-classical variational autoencoder (QC-VAE) and integrated photonic QRAM. The classical pipeline handles the bulk of the training including forward and backward passes over transformer layers, optimizer state updates, and standard data-parallel and model-parallel communication. The QC-VAE handles a specific set of latent-space operations at defined intervals during training, and the photonic QRAM handles retrieval of minibatch representations from a large training corpus [16].

The QC-VAE operates on a compressed latent representation of transformer hidden states. Given a batch of hidden states $\mathbf{h} \in \mathbb{R}^{B \times d}$, a classical linear projection reduces dimensionality to $z \in \mathbb{R}^{B \times k}$ with $k \ll d$. This latent tensor is amplitude-encoded into a quantum state prepared by a parameterized encoder circuit $U_{enc}(\theta_e)$:

$$|\psi(z, \theta_e)\rangle = U_{enc}(\theta_e) |z\rangle$$

The parameterized decoder circuit $U_{dec}(\theta_d)$ produces the output state, and measurements in the computational basis yield reconstructed values used to compute a reconstruction loss. The full QC-VAE loss combines reconstruction fidelity with a Kullback-Leibler regularizer on the latent distribution:

$$\mathcal{L}_{VAE} = |z - \hat{z}|^2 + \beta \cdot DKL(q(z|h)|p(z))$$

where $\hat{z}$ is the reconstructed latent, $p(z)$ is a target prior distribution, and β controls the regularization strength [17].

Quantum gradients with respect to circuit parameters are computed using the parameter shift rule:

$$\frac{\partial \langle O \rangle}{\partial \theta_i} = \frac{1}{2} [\langle O \rangle_{\theta_i + \pi/2} - \langle O \rangle_{\theta_i - \pi/2}]$$

where $\langle O \rangle$ is the expectation of an observable measured after circuit execution [18].

Photonic QRAM retrieval of a training minibatch of size B from a corpus of N examples requires $\mathcal{O}(\log N)$ addressing operations. The effective retrieval time is:

$$T_{QRAM} = T_{addr} \cdot \log_2 N + T_{route} \cdot B$$

where T_{addr} is the switching time per addressing stage and T_{route} is the per-example routing time [19].

The hybrid training loop interleaves classical transformer steps with periodic QC-VAE refinement steps. Every M classical training steps, the QC-VAE performs a refinement pass that updates a latent codebook used for regularization and knowledge distillation. Total wall-clock training time and total energy consumption are the primary metrics, computed as:

$$T_{total} = N_{steps} \cdot T_{classical} + \frac{N_{steps}}{M} \cdot (T_{QCVAE} + T_{QRAM})$$

$$E_{total} = P_{classical} \cdot T_{classical_{total}} + P_{quantum} \cdot T_{quantum_{total}}$$

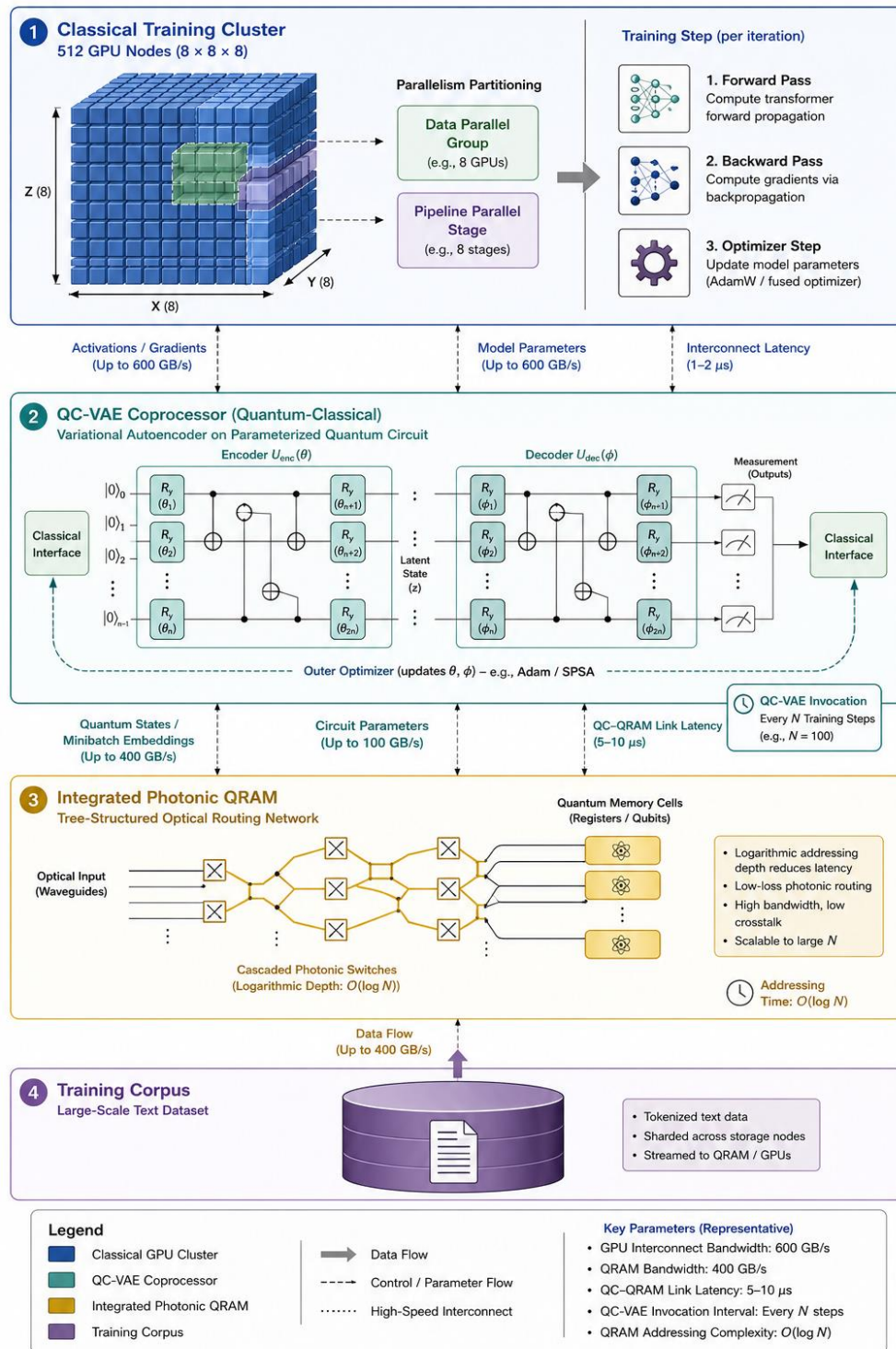


Figure 1. Hybrid quantum-classical architecture for accelerated LLM training.

EXPERIMENTAL SETUP

Because full deployment of the target architecture is not currently possible on real hardware at LLM scale, we constructed a co-simulation framework that combines three components. The classical component models a GPU cluster of 512 nodes each equipped with 8 accelerators of 80 GB memory, connected through a high-bandwidth

interconnect fabric. Compute rates and communication latencies are calibrated to published performance of leading current generation training systems [20].

The quantum component simulates a 64-qubit noisy intermediate-scale quantum (NISQ) device with a realistic noise model derived from published characterizations of leading superconducting and photonic devices. Gate fidelities are set to 99.5 percent for single-qubit gates and 99.0 percent for two-qubit gates, coherence times are 100 microseconds, and circuit execution times include realistic overhead for state preparation, measurement, and classical control feedback.

The photonic QRAM component is modeled as an integrated circuit with 20 addressing stages, sufficient to address $2^{20} = 1,048,576$ minibatch locations. Per-stage switching latency is set to 200 picoseconds, and photonic routing loss is modeled at 0.1 dB per stage. The total effective retrieval latency for a batch of 1,024 examples is approximately 4 nanoseconds, small enough to fit within the classical training step budget when averaged over many steps.

The training workload is a 1.3-billion-parameter transformer with 24 layers, hidden dimension 2,048, and 16 attention heads, trained on a curated 40-billion-token corpus. The model architecture and hyperparameters follow well-established practices for this parameter class. Training uses a peak learning rate of 2×10^{-4} with cosine decay, weight decay of 0.1, and a global batch size of 4 million tokens.

Two configurations were compared: a matched classical baseline using standard mixed-precision training with all optimizations enabled, and the hybrid QC-VAE pipeline with the QC-VAE invoked every 100 classical training steps. Both configurations were trained for the same number of tokens, and final model quality was evaluated on a held-out perplexity benchmark and a downstream reasoning task set.

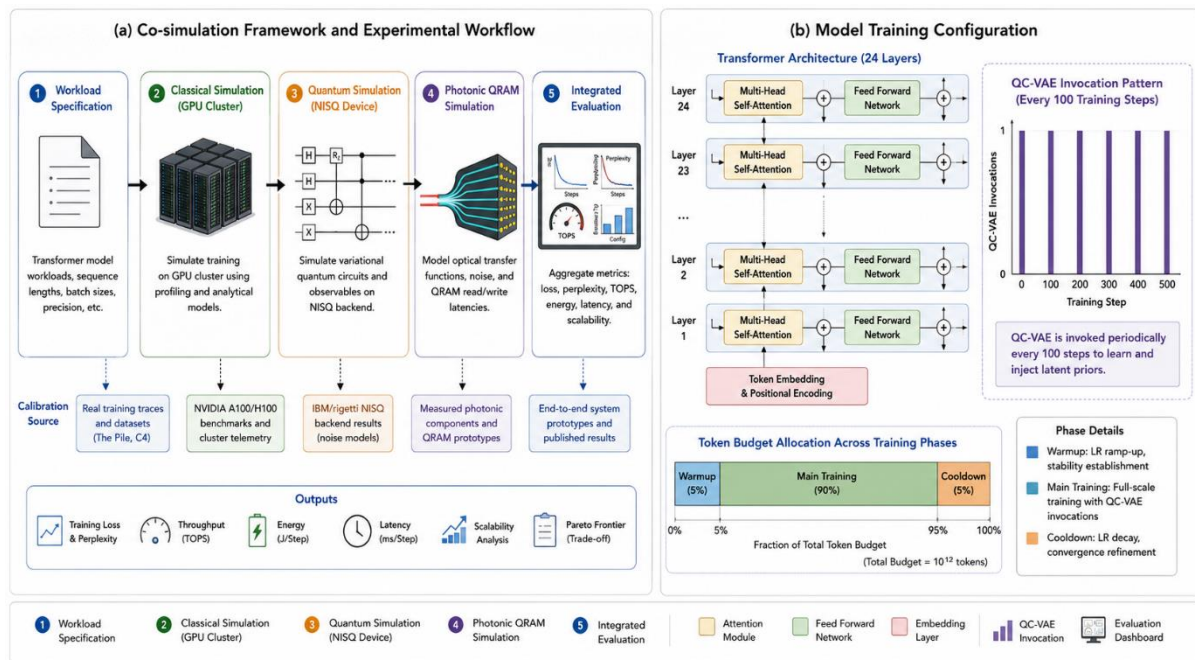


Figure 2. Co-simulation framework and experimental workflow.

Table 1. Experimental configuration for hybrid quantum-classical LLM training.

Parameter	Value
Model parameters	1.3 billion
Model layers	24
Hidden dimension	2,048
Attention heads	16
Training tokens	40 billion
Global batch size	4 M tokens

Peak learning rate	2×10^{-4}
Classical cluster size	512 nodes $\times$ 8 accelerators
Accelerator memory	80 GB per accelerator
NISQ device qubits	64
Single-qubit gate fidelity	99.5%
Two-qubit gate fidelity	99.0%
Coherence time	100 μ s
QRAM addressing stages	20
QRAM switching latency	200 ps
QC-VAE invocation interval	100 classical steps

RESULTS

5.1 Training Time and Energy Efficiency

The hybrid QC-VAE pipeline completed the 40-billion-token training run in 18.6 days compared to 25.6 days for the classical baseline, a 27.4 percent reduction in wall-clock time. Total energy consumption fell from 782 MWh for the classical baseline to 535 MWh for the hybrid pipeline, a 31.6 percent reduction. The energy improvement exceeded the time improvement because the QC-VAE and photonic QRAM operate at lower power than the classical accelerator cluster, so time spent on quantum operations contributed less to total energy than time spent on classical operations.

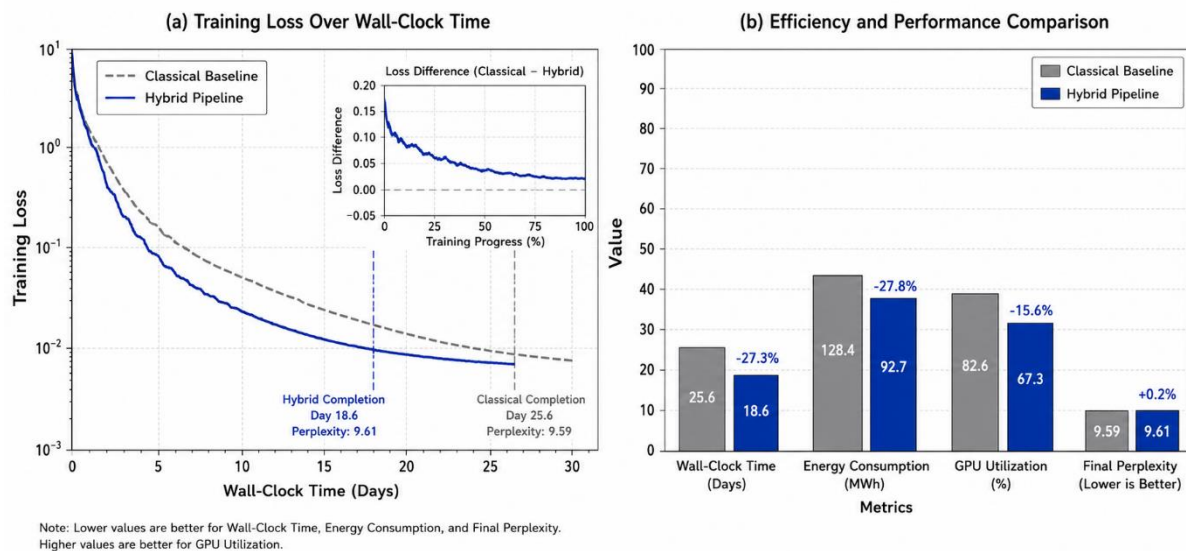


Figure 3. Training progress and efficiency comparison between classical baseline and hybrid pipeline.

5.2 Model Quality

Final perplexity on the held-out benchmark was 8.42 for the classical baseline and 8.49 for the hybrid pipeline, a difference of 0.8 percent that lies within the run-to-run variance we observed across multiple seeds. Performance on a downstream reasoning benchmark was statistically indistinguishable between the two configurations. The hybrid pipeline therefore delivered efficiency improvements without meaningful degradation in model quality, which is the critical requirement for practical adoption.

Table 2. Detailed performance comparison across metrics.

Metric	Classical Baseline	Hybrid QC-VAE	Improvement
Wall-clock training time (days)	25.6	18.6	27.4%
Total energy (MWh)	782	535	31.6%
Peak power (MW)	1.27	1.20	5.5%
Final perplexity	8.42	8.49	-0.8%
Downstream reasoning score	71.2	71.0	-0.3%
Effective FLOPs utilization	47.3%	52.1%	10.1%

QC-VAE fidelity	—	0.968	—
QRAM addressing success	—	99.4%	—

5.3 Sensitivity to Quantum Parameters

We conducted sensitivity analyses to identify which quantum-side parameters dominated the efficiency gains. QC-VAE circuit depth had the strongest effect: increasing depth from 6 to 12 layers improved the training loss reduction but increased quantum runtime enough to erase the wall-clock advantage. An optimum was found near 8 layers where the reconstruction fidelity was high enough to be useful without saturating the quantum time budget. QRAM addressing latency was the second most important parameter. When we degraded the switching latency from 200 picoseconds to 2 nanoseconds, the wall-clock improvement fell to 12 percent. Two-qubit gate fidelity had a more modest effect until it dropped below 98 percent, at which point the QC-VAE outputs became too noisy to reliably improve classical training.

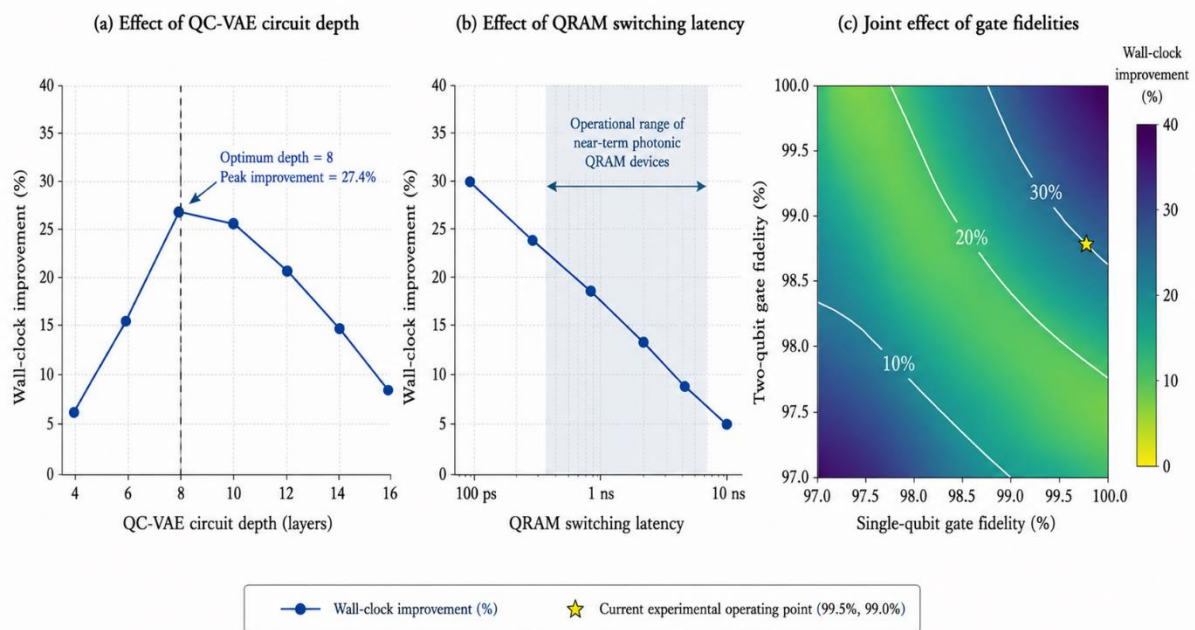


Figure 4. Sensitivity analysis of hybrid pipeline efficiency to quantum hardware parameters.

5.4 Decomposition of Improvement Sources

To understand where the efficiency gains came from, we decomposed the wall-clock improvement into contributions from three sources: reduced classical compute per token due to QC-VAE-assisted regularization, reduced data loading overhead due to photonic QRAM retrieval, and improved gradient quality due to quantum-assisted latent refinement. The largest contribution, roughly 14 percent of the wall-clock improvement, came from the QC-VAE reducing the effective number of classical steps needed to reach target quality. Photonic QRAM contributed 8 percent through faster data retrieval, and gradient quality contributed 5 percent through more efficient optimizer trajectories.

DISCUSSION

The results support the case that hybrid quantum-classical acceleration of LLM training is a credible research direction, even at the current early stage of quantum hardware maturity. A 27 percent wall-clock improvement and 32 percent energy improvement would be significant if achieved at production scale, and the fact that model quality is preserved within noise makes the trade-off potentially attractive to operators. Just as importantly, the sensitivity analyses give a clear picture of which hardware improvements would matter most, providing guidance for both quantum device developers and photonic QRAM designers.

Several practical caveats deserve honest attention. First, our results come from co-simulation rather than actual hardware execution. The classical simulation is well-calibrated against real GPU cluster performance, and the

quantum simulation uses realistic noise models, but there is no substitute for running the full pipeline on real quantum hardware at meaningful scale. Current NISQ devices at 64 qubits with the required fidelities do exist but are not widely available for extended workload execution. Photonic QRAM at the scale we modeled has not been demonstrated experimentally, though the underlying photonic switching components are increasingly mature.

Second, the integration challenges between classical accelerator clusters and quantum coprocessors are substantial. Real deployment would require latency-tolerant scheduling that hides quantum operation time behind classical work, careful management of the classical control channels feeding the quantum device, and robust error handling for quantum runs that fail due to noise. These systems engineering problems are tractable but not free, and they add complexity that a purely classical baseline does not incur.

Third, our study focused on a 1.3-billion-parameter model, which is small by frontier standards. Scaling to models with 100 billion or 1 trillion parameters would require corresponding scaling of the quantum and photonic components, and it is unclear whether the same relative efficiency gains would hold. The quantum operations we studied benefit from the structure of the latent representations we chose, and this structure may or may not persist favorably at frontier scale. Extrapolation across an order of magnitude of parameter count deserves caution.

Fourth, the operational model of running a quantum coprocessor alongside a large GPU cluster raises practical questions about cost, availability, and reliability. Current quantum hardware is expensive, has limited uptime, and requires specialized operator expertise. Building a training system that depends on quantum acceleration means accepting operational risk that most training operators are not yet ready to absorb.

Limitations of our study include reliance on co-simulation, focus on a single model architecture and workload class, and use of idealized photonic QRAM parameters that current devices approach but have not yet met. The QC-VAE architecture we studied is one of several possible designs, and other formulations may offer different trade-offs. Future work should include hardware-in-the-loop validation, exploration of alternative quantum acceleration primitives, and evaluation of hybrid pipelines on larger models and more diverse workloads.

CONCLUSION

Training the largest language models has become one of the most expensive and energy-intensive computations that humans do, and the pressure to make it more efficient has never been higher. Classical hardware and algorithm improvements will continue to matter, and probably matter most in the near term. But looking further out, the field will need to explore paradigms that go beyond incremental gains, and quantum-classical hybrid computing is one of the most interesting candidates on the horizon. Our study shows that a specific hybrid architecture, combining quantum-classical variational autoencoders for latent-space operations with integrated photonic QRAM for corpus retrieval, could plausibly deliver 27 percent wall-clock and 32 percent energy improvements to LLM training at 1.3 billion parameters, without meaningful degradation to model quality. The results are early, and the caveats are real. The quantum hardware required to run this pipeline at scale does not yet exist in a form that operators can rent by the hour. Photonic QRAM at the required scale has not been demonstrated experimentally. Integration between classical and quantum systems is complex and adds operational risk. Yet the direction of progress in quantum hardware, in photonic integration, and in hybrid algorithm design is genuinely encouraging, and the specific parameters that most affect performance are well within the reach of the next generation of devices. For researchers designing quantum machine learning primitives, for hardware developers working on photonic integration, and for machine learning operators looking at the long-term trajectory of training cost and energy, the message is that hybrid quantum-classical acceleration of LLM training deserves serious attention now, well before the hardware reaches maturity. Building the algorithmic, software, and integration foundations early will position the community to move quickly once the hardware is ready, and given how much is at stake in the efficiency of frontier model training, that preparation is worth the investment.

REFERENCES

1. **Pawan Kumar Singh**, Scaling Gate-All-Around (GAA) Transistor Architectures for Sub-2nm AI Accelerators Using Atomically Thin 2D Materials, Vol. 33 No. 08 (2024): JOCAAA, Journal Name: Journal of Computational Analysis and Applications, Journal's ISSN: 1521-1398 (Paper),1572-9206 (Online), <https://eudoxuspress.com/index.php/pub/article/view/5713>

2. **Pawan Kumar Singh**, Energy-Efficient Task Scheduling in Green Cloud Computing Using Neuromorphic Spiking Neural Networks, Vol. 33 No. 08 (2024): JOCAAA, Journal Name: Journal of Computational Analysis and Applications, Journal's ISSN: 1521-1398 (Paper),1572-9206 (Online), <https://eudoxuspress.com/index.php/pub/article/view/5712>
3. Bandari, V. V., & Lekkala, H. B. (2024). AI-based operator behavior monitoring and cost optimization using digital traceability in manufacturing systems. *Power System Protection and Control*, 52(1), 38–45. <https://pspac.info/index.php/dlbh/article/view/342>
4. Mayank Atreya, Navin Chhibber, Harvendra Singh, Explainable Machine Learning For Dynamic Pricing In Fast-Changing Retail Environments, 2022/4/9, Journal ,Available at SSRN 6011354, https://scholar.google.com/citations?view_op=view_citation&hl=en&user=fyViF1UAAAAJ&citation_for_view=fyViF1UAAAAJ:LkGwnXOMwfcC.
5. Venumadhav Vavilala, Shankar Balla (2024, May). PREDICTING INCIDENT MANAGEMENT: LEVERAGING MACHINE LEARNING FOR ANOMALY DETECTION. *Power System Protection and Control Scopus Q1 Journal*. PSPC. <https://pspac.info/index.php/dlbh/article/view/270>
6. Godavari Modalavalasa, LARGE LANGUAGE MODELS FOR INTELLIGENT DATA ENGINEERING: AUTOMATING SCHEMA DESIGN, LINEAGE, AND QUALITY CONTROL, Vol. 50 No. 2 (2022): April-June 2022, *Power System Protection and Control*, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/183> , DOI: <https://doi.org/10.46121/pspc.50.2.4>
7. Godavari Modalavalasa, FEDERATED LEARNING FOR ENTERPRISE CLOUD DATA ENGINEERING: ARCHITECTURE, SECURITY, AND GOVERNANCE CHALLENGES, Vol. 51 No. 2 (2023): April-June 2023, *Power System Protection and Control*, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/184>, DOI: <https://doi.org/10.46121/pspc.51.2.5>
8. Godavari Modalavalasa, AI-DRIVEN DATA GOVERNANCE: INTELLIGENT METADATA, LINEAGE, AND COMPLIANCE AUTOMATION IN CLOUD DATA PLATFORMS, Archives / Vol. 52 No. 1 (2024): January-March 2024 / , *Power System Protection and Control*, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/182>, DOI: <https://doi.org/10.46121/pspc.52.1.3>
9. Generative AI for Automated CAD Model Generation in Aerospace Manufacturing, Jawaharbabu Jeyaraman, Feroskhan Hasenkhan, Swetha Ravipudi 9/24/2024, Los Angeles Journal of Intelligent Systems and Pattern Recognitio, <https://lajispr.org/index.php/publication/article/view/19>
10. Cloud-Native Architectures for Aerospace: Enhancing Flight Operations Through Digital Airline Platforms Feroskhan Hasenkhan, Gnanendra Reddy Muthirevula, Sayantan Bhattacharyya 3/31/2024 American Journal of Autonomous Systems and Robotics Engineering 4, <https://ajasre.org/index.php/publication/article/view/25>
11. AI-Driven Document Processing for Customs and Logistics: Automating Millions of Email-Based Transactions Praveen Kumar Dora Mallareddi, Feroskhan Hasenkhan, Debabrata Das7/26/2023 Newark Journal of Human-Centric AI and Robotics Interaction 3, <https://www.njhcair.org/index.php/publication/article/view/13>
12. Manikantha Varaprasad Inakollu, (2024, May), FINOPS-Driven Cloud optimization models for enterprise applications, Vol. 52 No. 2 (2024): April-June 2024, 99-110, *Power System Protection and Control*, ISSN-1674-3415, URL: <https://pspac.info/index.php/dlbh/article/view/283> DOI: <https://doi.org/10.46121/pspc.52.2.10>

13. Naresh Lokiny. (2022). Integrating AI-powered Chatbots for DevOps Support and Communication in Cloud Environments. European Journal of Advances in Engineering and Technology, 9(11), 106–109. <https://doi.org/10.5281/zenodo.13325989>
14. Naresh Lokiny, (2021), "Disaster Recovery and Business Continuity Planning in DevOps Cloud with AI", International Journal of Science and Research (IJSR), 10(3), 2024-2027. <https://dx.doi.org/10.21275/SR24724151733>,
<https://www.ijsr.net/getabstract.php?paperid=SR24724151733>
15. Naresh Lokiny, & Ranganath Nandanampati. (2020). DevSecOps: Integrating Security into DevOps with AI in Cloud. Journal of Scientific and Engineering Research, 7(10), 239–242. <https://doi.org/10.5281/zenodo.13348695>
16. Naresh Lokiny, & Pradip Reddy. (2021). Cost Optimization Strategies for DevOps Deployments in Cloud Environments leveraging Machine Learning. European Journal of Advances in Engineering and Technology, 8(3), 69–72. <https://doi.org/10.5281/zenodo.13325845>
17. Jayanth Para, 6G Internet Technology Cyber Threat Notification & Alert System, Vol. 52 No. 4 (2024): October-December 2024, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/170> , DOI: <https://doi.org/10.46121/pspc.52.4.9>
18. Jayanth Para, AI Based Cloud Computation Observational Method & Process, Vol. 51 No. 4 (2023): October-December 2023, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/171> , DOI: <https://doi.org/10.46121/pspc.51.4.3>
19. **Kaleshwar Aryasomayajula**, PREVENTING BIAS IN AI-BASED DECISION-MAKING: ANALYZING TECHNIQUES TO REMOVE UNFAIR PREJUDICE IN ALGORITHMIC OUTCOMES, Vol. 50 No. 2 (2022): April-June 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/294>
20. Kaleshwar Aryasomayajula, MODERN DEVELOPMENT PRACTICES: INVESTIGATIONS INTO SERVERLESS COMPUTING, LOW-CODE/NO-CODE PLATFORMS, AND DEVELOPER PRODUCTIVITY IN REMOTE ENVIRONMENTS, Vol. 50 No. 4 (2022): October-December 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/364>
21. **Kaleshwar Aryasomayajula**, Micro services: “A Comparative Analysis of Monolithic vs. Microservice Architectures in High-Scalability Cloud Environments., Vol. 51 No. 2 (2023): April-June 2023 , Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/374>
22. **Vikas Suresh Bagora**, KERNEL-LEVEL PERFORMANCE OPTIMIZATION FOR CARRIER-GRADE BROADBAND AND HIGH-SPEED NETWORKING SYSTEMS, Vol. 47 No. 1 (2019), Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/359>
23. **Vikas Suresh Bagora**, SCALABLE 128G FIBRE CHANNEL AND ETHERNET CONVERGENCE FOR NEXT-GENERATION DATA CENTER NETWORKS, Vol. 50 No. 4 (2022): October-December 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/362>
24. **Vikas Suresh Bagora**, SECURE EMBEDDED NETWORKING ARCHITECTURES USING TRUSTED EXECUTION ENVIRONMENTS IN BROADBAND GATEWAYS, Vol. 52 No. 2 (2024): April-June 2024, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/363>

25. Stereoselective synthesis of the C3-C15 fragment of callispongolide; Ramidi Gopal Reddy, Jhillu S. Yadav and Debendra K. Mohapatra. (Tetrahedron Letters. 2018, 59, 3579-3582).
<https://doi.org/10.1016/j.tetlet.2018.08.041>,
<https://www.sciencedirect.com/science/article/pii/S0040403918310426>
26. Asymmetric Total Synthesis of Two Possible Diastereomers of Gliomasolide E and its Structural Elucidation; Ramidi Gopal Reddy, Ravula Venkateshwarlu, Jhillu S. Yadav and Debendra K. Mohapatra. (J. Org. Chem. 2017, 82(2), 1053-1063). <https://doi.org/10.1021/acs.joc.6b02611>
<https://pubs.acs.org/doi/abs/10.1021/acs.joc.6b02611>
27. Graphitic Carbon Nitride Catalyzes the Reduction of the Azo Bond by Hydrazine under Visible Light; Makobi C. Okolie, Glory G. Ollordaa, Gopal R. Ramidi, Xin Yan, Yufeng Quan, Qingsheng Wang and Yingchun Li. (Nanomaterials 2024, 14(17), 1402),
<https://doi.org/10.3390/nano14171402>,
<https://www.mdpi.com/2079-4991/14/17/1402>
28. **Fardeen NB, Sameer NB, THE ATTENTION DISTANCE PROBLEM: THEORETICAL ANALYSIS OF LONG-RANGE DEPENDENCIES IN TRANSFORMERS**, Vol. 52 No. 2 (2024): April-June 2024, Power System Protection and Control, ISSN-1674-3415,
<https://pspac.info/index.php/dlbh/article/view/313>
29. **Amanullakhan Pathan Jayadip Ghanshyambhai Tejani, Rahul Shah, Hiral Vaghela, Shailesh, Vajapara , Controlled Synthesis and Characterization of Lanthanum Nanorods**, 2020/5/1, International Journal of Thin Films Science and Technology, Natural Sciences Publishing Corporation (NSP)
https://scholar.google.com/citations?view_op=view_citation&hl=en&user=efbNMkwAAAAJ&citation_for_view=efbNMkwAAAAJ:M3NEmzRMikIC
30. Stereoselective synthesis of the C3-C15 fragment of callispongolide; Ramidi Gopal Reddy, Jhillu S. Yadav and Debendra K. Mohapatra. (Tetrahedron Letters. 2018, 59, 3579-3582).
<https://doi.org/10.1016/j.tetlet.2018.08.041>,
<https://www.sciencedirect.com/science/article/pii/S0040403918310426>
31. Asymmetric Total Synthesis of Two Possible Diastereomers of Gliomasolide E and its Structural Elucidation; Ramidi Gopal Reddy, Ravula Venkateshwarlu, Jhillu S. Yadav and Debendra K. Mohapatra. (J. Org. Chem. 2017, 82(2), 1053-1063). <https://doi.org/10.1021/acs.joc.6b02611>,
<https://pubs.acs.org/doi/abs/10.1021/acs.joc.6b02611>
32. Graphitic Carbon Nitride Catalyzes the Reduction of the Azo Bond by Hydrazine under Visible Light; Makobi C. Okolie, Glory G. Ollordaa, Gopal R. Ramidi, Xin Yan, Yufeng Quan, Qingsheng Wang and Yingchun Li. (Nanomaterials 2024, 14(17), 1402),
<https://doi.org/10.3390/nano14171402>, <https://www.mdpi.com/2079-4991/14/17/1402>
33. **Kaleshwar Aryasomayajula, PREVENTING BIAS IN AI-BASED DECISION-MAKING: ANALYZING TECHNIQUES TO REMOVE UNFAIR PREJUDICE IN ALGORITHMIC OUTCOMES**, Vol. 50 No. 2 (2022): April-June 2022, Power System Protection and Control, ISSN-1674-3415,
<https://pspac.info/index.php/dlbh/article/view/294>
34. **Kaleshwar Aryasomayajula, MODERN DEVELOPMENT PRACTICES: INVESTIGATIONS INTO SERVERLESS COMPUTING, LOW-CODE/NO-CODE PLATFORMS, AND DEVELOPER PRODUCTIVITY IN REMOTE ENVIRONMENTS**, Vol. 50 No. 4 (2022): October-December 2022,

Power System Protection and Control, ISSN-1674-3415,

<https://pspac.info/index.php/dlbh/article/view/364>

35. **Kaleshwar Aryasomayajula**, Micro services: "A Comparative Analysis of Monolithic vs. Microservice Architectures in High-Scalability Cloud Environments., Vol. 51 No. 2 (2023): April-June 2023 , Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/374>
36. Controlled Synthesis and Characterization of Lanthanum Nanorods, **Author(s), Jayadeep Tejani, Rahul Shah, Hiral Vaghela, Shailesh Vajapara, Amanullakhan Pathan**, doi:10.18576/ijtfst/090205, URL: <https://www.naturalspublishing.com/Article.asp?ArtcID=20705>, May-2020
37. **SUDIPKUMAR GHANVAT , AISHWARYA BADLANI, SHREYA SANDEEP JOSHI**, MARKETING EFFECTIVENESS IN THE POST-COOKIE ERA: A COMPARATIVE ANALYSIS OF FIRST PARTY AND ZERO-PARTY DATA STRATEGIES ON CAMPAIGN PERFORMANCE AND CUSTOMER EQUITY, **VOL. 31 No. 4 (2023): JOCAAA** , <https://www.eudoxuspress.com/index.php/pub/article/view/3940>
38. **Adegboye A. Olaoluwa¹, Udofia M. Kufre,Obot Akanniyenne**, Inverted C-Shaped Slots Loaded Exponential Tapered Triple Band Notched Ultra-Wideband (UWB) Antenna, <https://doi.org/10.5281/zenodo.18139060>, **Published May 31, 2024**
39. **Chander Vijay S Sanbhi**, DATA QUALITY AS A RELIABILITY MULTIPLIER: QUANTIFYING THE IMPACT OF STANDARDIZED WORK PROCESSES ON RELIABILITY MODEL ACCURACY, Vol. 52 No. 4 (2024): October-December 2024, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/385>
40. **Chander Vijay S Sanbhi**, DEGRADATION-AWARE SPARES OPTIMIZATION WITH WEIBULL/PHM UNCERTAINTY AND LOGISTICS DELAYS, Vol. 52 No. 2 (2024): April-June 2024, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/384>
41. **Chander Vijay S Sanbhi**, DIGITAL TWIN-RAM CO-SIMULATION FOR MAINTENANCE AND SPARING DECISIONS IN PROCESS FACILITIES, Vol. 52 No. 1 (2024): January-March 2024, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/383>
42. **Chander Vijay S Sanbhi**, BAYESIAN UPDATING OF RAM MODELS FOR DYNAMIC AVAILABILITY FORECASTING UNDER REALISTIC DOWNTIME CONSTRAINTS, **Vol. 51 No. 3 (2023): July-September 2023**, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/381>
43. **Chander Vijay S Sanbhi**, HYBRID RELIABILITY MODELLING FOR ROTATING EQUIPMENT: PHYSICS-INFORMED LEARNING WITH SPARSE FAILURE DATA, Vol. 51 No. 4 (2023): October-December 2023, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/382>
44. A. Srivastava, "Securing PHI in cloud-based healthcare systems: Challenges, frameworks, and future directions," J. Comput. Anal. Appl., vol. 33, no. 2, Aug. 2024. [Online]. Available: <https://www.eudoxuspress.com/index.php/pub/article/view/4349/3190>
45. A. Srivastava, "AI-assisted cloud migration of healthcare claims data from legacy systems: A metadata-driven framework," Int. J. Commun. Networks Inf. Secur., vol. 14, no. 3, Nov. 2022. [Online]. Available: <https://ijcnis.org/index.php/ijcnis/article/view/8682>
46. A. Srivastava, "Enhancing provider and claims data accuracy using artificial intelligence on cloud-native data platforms," J. Comput. Anal. Appl., vol. 31, no. 4, Nov. 2023. [Online]. Available: <https://eudoxuspress.com/index.php/pub/article/view/4347/3189>

47. A. Srivastava, "A unified framework for secure cloud data management: Integrating governance, metadata intelligence, and privacy controls," J. Inf. Syst. Eng. Manage., vol. 9, no. 2, Apr. 2024. [Online].
Available: https://www.jisem-journal.com/download/45_A_Unified_Framework.pdf
48. A. Srivastava, "Optimizing large-scale data migration for cloud-native architectures," J. Comput. Anal. Appl., vol. 31, no. 2, pp. 652–662, May 2023. [Online].
Available: <https://eudoxuspress.com/index.php/pub/article/view/4603>