

# Quantum-Enhanced Attention Mechanism in NLP: A Hybrid Classical-Quantum Approach

by

S.M. Yousuf Iqbal Tomal

21301129

Abdullah Al Shafin

21201631

Debojit Bhattacharjee

20201159

MD. Khairul Amin

21201167

A thesis submitted to the Department of Computer Science and Engineering  
in partial fulfillment of the requirements for the degree of  
B.Sc. in Computer Science and Engineering

Department of Computer Science and Engineering  
Brac University  
June 2025

© 2024. Brac University  
All rights reserved.

# Declaration

It is hereby declared that

1. The thesis submitted is our own original work while completing degree at Brac University.
2. The thesis does not contain material previously published or written by a third party, except where this is appropriately cited through full and accurate referencing.
3. The thesis does not contain material which has been accepted, or submitted, for any other degree or diploma at a university or other institution.
4. We have acknowledged all main sources of help.

**Student's Full Name and Signature:**

---

S.M. Yousuf Iqbal Tomal  
21301129

---

Abdullah Al Shafin  
21201631

---

Debojit Bhattacharjee  
20201159

---

MD. Khairul Amin  
21201167

# Approval

The thesis/project titled “Quantum-Enhanced Attention Mechanism in NLP: A Hybrid Classical-Quantum Approach” submitted by

1. S.M. Yousuf Iqbal Tomal (21301129)
2. Abdullah Al Shafin (21201631)
3. Debojit Bhattacharjee (20201159)
4. MD. Khairul Amin (21201167)

of Spring, 2025 has been accepted as satisfactory in partial fulfillment of the requirement for the degree of B.Sc. in Computer Science in June 20, 2025.

## Examining Committee:

Supervisor:  
(Member)

---

Rafiad Sadat Shahir

Lecturer  
Department of Computer Science and Engineering  
BRAC University

Thesis Coordinator:  
(Member)

---

Md. Golam Rabiul Alam, PhD

Professor  
Department of Computer Science and Engineering  
BRAC University

Head of Department:  
(Chair)

---

Sadia Hamid Kazi

Associate Professor  
Department of Computer Science and Engineering  
BRAC University

# Abstract

Recent advances in quantum computing have opened new pathways for enhancing deep learning architectures, particularly in domains characterized by high-dimensional and context-rich data such as natural language processing (NLP). In this work, we present a hybrid classical-quantum Transformer model that integrates a quantum-enhanced attention mechanism into the standard classical architecture. By embedding token representations into a quantum Hilbert space via parameterized variational circuits and exploiting entanglement-aware kernel similarities, the model captures complex semantic relationships beyond the reach of conventional dot-product attention. We demonstrate the effectiveness of this approach across diverse NLP benchmarks, showing improvements in both efficiency and representational capacity. Empirical study reveals that the quantum attention layer yields globally coherent attention maps and more separable latent features, while requiring comparatively fewer parameters than classical counterparts. These findings highlight the potential of quantum-classical hybrid models to serve as a powerful and resource-efficient alternative to existing attention mechanisms in NLP.

**Keywords:** Quantum Attention, Variational Quantum Circuit(VQC), Quantum Kernel, Hybrid Quantum-Classical Model, Natural Language Processing(NLP)

# Table of Contents

|                                                               |           |
|---------------------------------------------------------------|-----------|
| Declaration . . . . .                                         | 1         |
| <b>1 Introduction</b>                                         | <b>8</b>  |
| 1.1 Research Problem . . . . .                                | 9         |
| 1.2 Research Contributions . . . . .                          | 9         |
| <b>2 Literature Review</b>                                    | <b>11</b> |
| <b>3 Theoretical Study</b>                                    | <b>13</b> |
| 3.1 Architecture Overview . . . . .                           | 13        |
| 3.2 Quantum Computing Principles . . . . .                    | 15        |
| 3.3 Quantum Attention Mechanism . . . . .                     | 16        |
| 3.3.1 Quantum Kernel . . . . .                                | 17        |
| 3.3.2 Quantum Interference Attention . . . . .                | 18        |
| 3.3.3 Quantum Multi-Head Attention . . . . .                  | 19        |
| 3.3.4 Quantum Superposition Layer . . . . .                   | 21        |
| 3.3.5 Full Quantum Advantage Transformer . . . . .            | 22        |
| 3.3.6 Comparison with Classical Baseline . . . . .            | 23        |
| <b>4 Empirical Study</b>                                      | <b>26</b> |
| 4.1 Performance Comparison . . . . .                          | 26        |
| 4.2 Quantum Advantage Analysis . . . . .                      | 29        |
| 4.3 Attention and Feature Analysis . . . . .                  | 30        |
| 4.4 Performance Discussion . . . . .                          | 33        |
| 4.4.1 Macro-level Trends . . . . .                            | 33        |
| 4.4.2 Linking Theory to Practice . . . . .                    | 33        |
| 4.4.3 Dataset-Specific Insights . . . . .                     | 34        |
| 4.4.4 Training Dynamics and Efficiency . . . . .              | 34        |
| <b>5 Discussion</b>                                           | <b>35</b> |
| 5.1 Synthesis of Theoretical and Empirical Evidence . . . . . | 35        |
| 5.2 Interpretability and Representational Geometry . . . . .  | 35        |
| 5.3 Entangling Ansatz Architecture . . . . .                  | 35        |
| 5.4 Comparative Perspective . . . . .                         | 36        |
| 5.5 Limitations . . . . .                                     | 36        |
| 5.6 Broader Impact and Concluding Remarks . . . . .           | 36        |
| <b>6 Future Directions</b>                                    | <b>37</b> |
| 6.1 Optimising the Circuit Ansatz . . . . .                   | 37        |
| 6.2 Expanding Network Complexity . . . . .                    | 37        |
| 6.3 Task-Specific Ansatz and Network Design . . . . .         | 37        |

|                                           |           |
|-------------------------------------------|-----------|
| 6.4 Data-Centric Quantum Tuning . . . . . | 38        |
| <b>7 Conclusion</b>                       | <b>39</b> |

# List of Figures

|     |                                                                                                                                                                                                                                 |    |
|-----|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 3.1 | Flow diagram of the Quantum Advantage Transformer (QAT) . . . . .                                                                                                                                                               | 15 |
| 3.2 | Quantum Kernel Circuit. The circuit demonstrates the combination of parameterized rotations ( $RX, RZ, RY$ ) and the Controlled-NOT (CNOT) gate used to compute quantum token similarities. . . . .                             | 18 |
| 3.3 | Quantum Superposition Layer. The circuit demonstrates the combination of Hadamard gate, parameterized rotations ( $RY, RZ$ ) and the Controlled-NOT (CNOT) gate. . . . .                                                        | 22 |
| 3.4 | Classical Transformer Encoder Architecture (used as baseline) . . . . .                                                                                                                                                         | 24 |
| 4.1 | Training loss trajectories (left of each panel) and final test accuracy (right) for QET Single-Head(blue), QET Multi-Head(purple), the classical multi-head Transformer (orange), and its single-head ablation (green). . . . . | 28 |
| 4.2 | Token-token interaction maps for QET (left), classical multi-head (centre) and classical single-head (right). Darker colours denote higher weights; colour bars are scaled independently per subplot. . . . .                   | 31 |
| 4.3 | t-SNE projections of QET sentence embeddings (colour = gold label). Compared to classical maps (Appendix C), clusters are tighter and inter-class gaps wider, indicating improved feature separation. . . . .                   | 32 |

# List of Tables

|     |                                                  |    |
|-----|--------------------------------------------------|----|
| 3.1 | QET Architectural Parameters . . . . .           | 14 |
| 4.1 | Model Parameter Counts . . . . .                 | 26 |
| 4.2 | Performance Comparison Across Datasets . . . . . | 28 |
| 4.3 | Quantum Advantage Cases (Multi-Head) . . . . .   | 29 |
| 4.4 | Quantum Advantage Cases . . . . .                | 30 |

# Chapter 1

## Introduction

Natural Language Processing (NLP) has rapidly become a cornerstone of modern artificial intelligence, enabling machines to interpret, generate, and interact with human language across a myriad of applications from machine translation and sentiment analysis to conversational agents and document summarization. These applications have transformed industries such as healthcare, finance, and customer service by allowing for more natural and efficient communication between humans and machines.

At the heart of recent breakthroughs in NLP lie transformer-based architectures such as BERT and GPT. These models employ self-attention mechanisms to adeptly capture long-range dependencies and intricate contextual cues across entire sequences. This architectural shift has yielded significant improvements in performance across a wide range of language understanding tasks. However, the success of these models has come at a cost: they demand increasingly large computational resources and memory, presenting challenges for their deployment in real-time systems or on devices with limited processing capabilities.

In parallel, quantum computing has emerged as a promising paradigm that offers fundamentally new ways of performing computations. By leveraging the principles of superposition and entanglement, quantum computers can encode and manipulate information in ways that are not feasible on classical machines. Quantum machine learning (QML), an interdisciplinary field at the intersection of quantum computing and artificial intelligence, aims to harness these capabilities to develop more efficient and expressive learning algorithms. In the subdomain of Quantum Natural Language Processing (QNLP), recent surveys such as Widdows et al. (2024) provide a comprehensive overview of how quantum principles are being applied to linguistic tasks, outlining both theoretical foundations and emerging practical implementations. [8]. Specifically, variational quantum circuits (VQCs) and quantum kernel methods allow data to be encoded into high-dimensional Hilbert spaces, potentially offering exponential representational power with relatively few parameters, while recent work such as Gao et al. (2023) demonstrates how quantum algorithms can accelerate attention computation [12].

Despite this promise, purely quantum systems face practical limitations. Current quantum hardware is constrained by noise, limited qubit counts, and challenges in scalability. These barriers make it difficult for standalone quantum models to handle the complexity and size of real-world NLP datasets. To bridge this gap, hybrid quantum-classical architectures have emerged as a compelling alternative which was shown by a workflow by O’Riordan et al. (2020) [15]. These models integrate quantum components into classical frameworks, strategically offloading the most computationally intensive tasks to quantum

processors while relying on classical modules for data pre-processing, model training, and inference which Dey et al. (2024) had shown in his recent work on text classification using the hybrid-quantum approach [10].

In this thesis, we propose the Quantum-Enhanced Transformer (QET), a hybrid model that augments standard self-attention mechanisms with quantum feature extraction and phase-controlled interference. Our approach integrates variational quantum circuits into the attention mechanism of transformer models, allowing for the capture of non-linear and higher-order dependencies between tokens. By combining quantum modules with the proven residual and normalization structures of classical transformer encoders, QET aims to enhance performance while reducing parameter complexity.

We empirically evaluate QET on several benchmark NLP datasets, including AG News, IMDB, SST-2, and SST-5. Our results demonstrate that QET achieves comparable or superior performance to classical baselines, while operating with fewer parameters. Notably, quantum attention is shown to enhance the separability of embedding spaces, distribute attention more evenly across tokens, and widen decision margins. These findings suggest that quantum inductive biases can be harnessed to develop more efficient and expressive NLP models, paving the way for the next generation of quantum-classical hybrid systems in language understanding.

## 1.1 Research Problem

Transformer-based models have revolutionized NLP, but their computational intensity poses scalability challenges, especially for large datasets and real-time applications. Quantum computing offers potential solutions through its unique properties, such as superposition and entanglement, which could enhance computational efficiency. Despite this potential, significant challenges remain in effectively integrating quantum computing into classical NLP architectures.

These challenges include managing the high error rates in current quantum computations, optimizing quantum models for handling large datasets, and ensuring interpretability in hybrid systems. Furthermore, there is a lack of established methods for encoding NLP data into quantum states and for optimizing hybrid quantum-classical models.

The specific problem addressed by this research is: Can a hybrid classical-quantum transformer model employing quantum-enhanced attention improve computational efficiency for NLP tasks without sacrificing accuracy? This study seeks to determine whether the integration of quantum kernel similarity and variational quantum circuits can reduce the computational burden of attention mechanisms in transformers, making NLP models more efficient and scalable.

## 1.2 Research Contributions

This work delivers the following contributions to hybrid quantum-classical NLP, merging our research objectives with achieved outcomes:

1. **Quantum-Enhanced Transformer (QET).** We designed a drop-in hybrid attention module that embeds token embeddings into an entangled Hilbert space via a two-layer, six-qubit variational circuit, blending quantum kernel similarity with classical dot-product attention. This achieves  $\approx 5\%$  parameter savings and

comparable accuracy gains on topic classification (AG NEWS) versus classical size-matched Transformers.

2. **Phase-Diverse Multi-Head Attention.** We extended our design to a multi-head regime with head-specific phase parameters that diversify interference patterns while maintaining drop-in compatibility. The four-head variant widens kernel expressivity, adding 0.28 % on long-form sentiment (IMDB) with less than half the parameter overhead of classical multi-head upgrades.
3. **Interpretability Toolkit.** Through empirical evaluation across sentiment analysis benchmarks (SST-2, SST-5), we developed spectral and t-SNE diagnostics that reveal smoother attention bands and tighter class clusters, providing geometric evidence that entanglement reshapes NLP representations.
4. **Hardware-Aware Circuit Ansatz.** By analyzing resource efficiency and hardware realism, we established that quantum gains outweigh circuit overhead using a reusable two-layer, six-qubit strongly-entangling circuit that balances expressivity, optimization stability, and NISQ feasibility.

Together, these contributions advance quantum-classical integration for language understanding, demonstrating measurable benefits in accuracy, efficiency, and interpretability while providing a blueprint for scalable hybrid architectures.

# Chapter 2

## Literature Review

Quantum computing, grounded in principles such as superposition and entanglement, has gained significant traction in the machine learning community due to its potential to outperform classical algorithms in select domains. Quantum Machine Learning (QML) integrates these quantum phenomena into traditional learning pipelines, offering potential improvements in both computational efficiency and representational power. As a foundation for our proposed quantum-enhanced attention mechanism, this section reviews key developments in classical NLP, quantum computing, and emerging hybrid quantum-classical models.

A pivotal advancement in NLP came with the introduction of the transformer architecture by Vaswani et al. (2017), which replaced recurrence and convolution with self-attention mechanisms [1]. This shift enabled significantly more parallelism and led to improved results across a wide range of NLP benchmarks. Building upon this, BERT (Devlin et al., 2019) introduced bidirectional attention, further enhancing performance on tasks such as sentiment analysis and question answering [2]. However, as model size and complexity grow, these classical systems face diminishing returns, particularly in terms of memory footprint, compute cost, and feasibility in real-time or resource-constrained environments.

To address these limitations, researchers have increasingly turned to quantum computing. Schuld et al. (2015), in *An Introduction to Quantum Machine Learning*, suggest that quantum encoding and quantum algorithms could offer NLP models new tools for handling high-dimensional data more efficiently [3]. Among the most promising directions are quantum kernel methods and variational quantum circuits (VQCs). These techniques offer a compact means of encoding complex semantic relationships into high-dimensional Hilbert spaces, potentially surpassing the limitations of classical attention mechanisms.

Cerezo et al. (2021) presented a formal framework for VQCs, demonstrating how they can enhance model expressiveness and optimize parameter space traversal [4]. While these circuits have shown success in broader machine learning contexts, their targeted integration into transformer attention modules remains underexplored which is a key motivation behind our work.

Important foundational contributions by Biamonte et al. (2017) highlight how quantum algorithms such as Grover’s search and quantum support vector machines (QSVMs) can improve feature extraction and reduce computational complexity [5]. These theoretical insights lay the groundwork for adapting quantum operations to NLP, particularly in enhancing self-attention modules, though practical implementations in this domain remain limited.

Further, Schuld et al. (2019) show that quantum kernel methods can capture non-linear, high-dimensional data distributions more effectively than classical techniques [6]. Their work suggests that quantum-enhanced representations could help NLP systems mitigate the curse of dimensionality, even though applications to attention remain scarce.

In 2020, Schuld et al. (2020) also introduced circuit-centric quantum classifiers, broadening the design space for integrating quantum circuits directly into classification pipelines [17]. Complementary to this, Mari et al. (2020) explored hybrid transfer learning, showing that pre-trained classical models can benefit from fine-tuning via quantum layers, improving both generalization and sample efficiency [9]. While these works do not directly tackle attention mechanisms, they reinforce the idea that classical–quantum integration can yield synergistic performance gains.

The advent of Quantum Convolutional Neural Networks (QCNNs) by Cong et al. (2019) also contributes to this landscape [11]. These networks exploit quantum parallelism to improve feature extraction in signal and image processing. While primarily applied outside of NLP, their architectural ideas especially in encoding locality and hierarchy suggest that quantum modules could meaningfully extend attention mechanisms in text-based models.

One notable attention-focused application is Quixer by Khatri et al. (2021), a quantum transformer that integrates quantum circuits into self-attention layers [7]. Their results show that quantum-enhanced architectures can better process complex attention patterns. However, practical barriers such as circuit depth, noise, and hardware scalability remain. Our proposed hybrid model addresses these concerns through lightweight quantum modules with minimal parameter overhead.

In another domain, Quantum LSTM (QLSTM) introduced by Chen et al. (2022) shows how quantum circuits can be embedded into recurrent structures to reduce parameter count without sacrificing the ability to capture sequential dependencies [13]. Although their focus is on RNNs, the architectural insight is broadly applicable to transformer designs as well.

Li, Zhao, and Wang (2024) further build on this trend with Quantum Self-Attention Neural Networks (QSANN), which embed word vectors directly into quantum states [14]. Their work shows improved classification performance, validating the use of quantum-enhanced attention mechanisms. However, their exploration stops short of kernel similarity modeling and VQC-based attention, which is the focal point of our proposed method.

Contributions from other fields also highlight the broad potential of quantum computing. For instance, Dong et al. (2008) explore quantum reinforcement learning, while Cherrat et al. (2022) present Quantum Vision Transformers for image tasks [16] [18]. These studies confirm the wide applicability of quantum methods and underscore the need for domain-specific adaptation in NLP.

In conclusion, while the body of research on quantum-enhanced models continues to grow, there remains a notable gap in the application of quantum principles to transformer attention mechanisms. This thesis aims to fill that gap by proposing a hybrid architecture that explicitly integrates VQCs and quantum kernel similarity into the self-attention layer, with the goal of enhancing both the efficiency and representational capacity of transformers for NLP tasks.

# Chapter 3

## Theoretical Study

In this section, we develop the formal framework underlying our Quantum-Enhanced Transformer (QET), a hybrid architecture that augments standard self-attention with quantum feature extraction. First, we define the QuantumKernel, which is a variational multi-qubit circuit that encodes the reduced embedding of each token through data-driven rotations and entangling gates, resulting in a high-dimensional quantum feature map. We then introduce an interference-based attention mechanism in which quantum feature inner products are adjusted by a learnable phase term before softmax normalization. Finally, we describe how these attention weights are reintegrated into the classical Transformer encoder to produce final token representations and classification logits. By exploiting quantum superposition and entanglement, QET projects token embeddings into a richer Hilbert space, enabling the model to capture long-range and nonlinear dependencies more effectively than purely classical counterparts, an improvement we demonstrate empirically in Section 4.

### 3.1 Architecture Overview

The Quantum-Enhanced Transformer (QET) architecture integrates classical transformer components with quantum-enhanced attention mechanisms to effectively model token dependencies and generate sentiment predictions. The workflow comprises the following stages:

- **Token Embedding Layer:** The input text is first tokenized into word-piece IDs using a pre-trained BERT tokenizer, producing sequences of fixed length. These IDs are then converted to continuous vectors via a learnable embedding layer. The resulting tensor captures contextualized token semantics. To interface with the quantum module, each embedding is linearly projected down to the quantum register size using a tensor. These vectors serve as input to the QuantumKernel circuit, ensuring a seamless transition from the classical embedding space to the quantum feature map without losing learned semantic information.
- **Quantum Attention Mechanism:** Reduced token embeddings are processed by a quantum-enhanced attention module that first encodes each vector into a multi-qubit state via a parameterized variational circuit, interleaving data-driven rotations and entangling gates to exploit superposition and entanglement. Pauli-Z measurements on each qubit yield a compact quantum feature vector for every token. Pairwise inner products of these quantum features form a similarity matrix,

to which a learnable phase-based interference term is added, before applying a row-wise softmax to produce attention weights. Finally, these weights scale the original embeddings, infusing classical self-attention with rich, non-linear quantum correlations while preserving the transformer’s residual and normalization pathways.

- **Classification Layer:** The attention-augmented token representations are aggregated into a fixed-size sequence summary via global mean pooling, producing a single feature vector per example. This pooled vector is then passed through a lightweight feedforward network comprising a linear layer that expands the feature dimensionality, a ReLU activation, dropout for regularization, and a final linear projection to the number of output classes (two for binary sentiment). The resulting logits are normalized with softmax during training to compute cross-entropy loss, enabling end-to-end learning of both classical and quantum parameters in a unified optimization.

**Data Flow:** The architecture operates in the following steps:

- **Tokenization & Embedding:** Text  $\rightarrow$  Token IDs  $\rightarrow$  Embeddings
- **Projection:** Embeddings  $\rightarrow$  Vectors for the quantum circuit.
- **Quantum Mapping:** Each vector  $\rightarrow$  Variational circuit  $\rightarrow$  Pauli-Z measurements  $\rightarrow$  Quantum features.
- **Attention Weights:** Compute pairwise dot products of quantum features, add a learnable interference term, and then softmax to obtain attention scores.
- **Re-weight & Superpose:** Apply scores to the original embeddings, then pass through a second variational circuit for extra entanglement.
- **Finalize:** Add residual connection, apply layer normalization, mean pool between tokens, and feed into a two-layer MLP to produce the final sentiment logits.

**Parameter Configuration:** The QET architecture is configured with the following key parameters:

Table 3.1: QET Architectural Parameters

| Parameter                  | Value                                                                                                       |
|----------------------------|-------------------------------------------------------------------------------------------------------------|
| Embedding Size             | 64                                                                                                          |
| Sequence Length            | IMDB: 120<br>AG NEWS: 120<br>SST-2: 64<br>SST-5: 120                                                        |
| Number of Qubits (VQC)     | 6                                                                                                           |
| Attention Heads            | Single-head: 1<br>Multi-head: 2                                                                             |
| Classification Output Size | IMDB: 2 (positive, negative)<br>AG NEWS: 4 (classes)<br>SST-2: 2 (positive, negative)<br>SST-5: 5 (classes) |

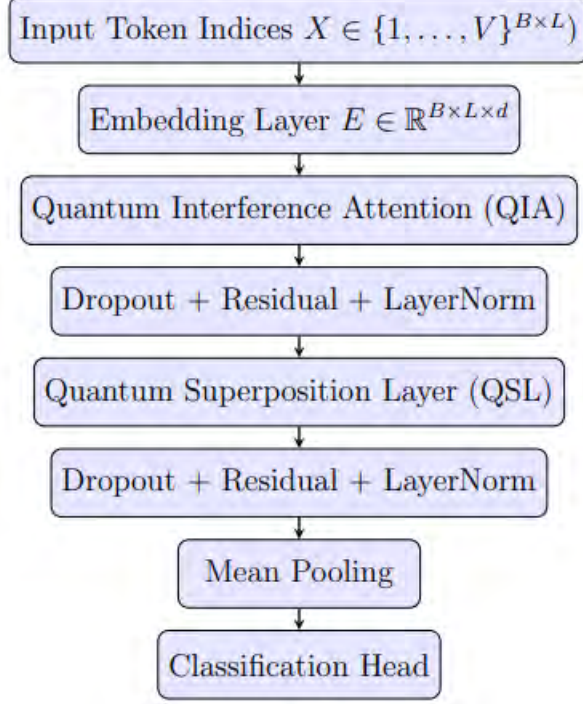


Figure 3.1: Flow diagram of the Quantum Advantage Transformer (QAT)

## 3.2 Quantum Computing Principles

Modern quantum computing leverages several foundational principles that distinguish it from classical paradigms. In this section, we succinctly outline these principles and their relevance to the proposed Quantum-Enhanced Transformer (QET).

**Quantum Superposition:** Quantum superposition permits a quantum register to occupy a linear combination of basis states simultaneously. Formally, for an  $n$ -qubit register, the state vector can be expressed as:

$$|\psi\rangle = \sum_{i=0}^{2^n-1} \alpha_i |i\rangle, \quad \sum_{i=0}^{2^n-1} |\alpha_i|^2 = 1, \quad (3.1)$$

where each amplitude  $\alpha_i$  encodes the probability amplitude of the basis state  $|i\rangle$ . This exponentially large Hilbert space, of dimension  $2^n$ , enables parallel encoding of classical data with only  $n$  qubits, offering a potential advantage in feature representation and exploration of solution spaces. In the context of the Quantum-Enhanced Transformer (QET), superposition allows token embeddings to be projected into a high-dimensional quantum feature space, facilitating richer encoding of semantic relationships and improved expressivity in attention computations.

**Quantum Entanglement:** Quantum entanglement generates intrinsic correlations between qubits such that the joint state cannot be decomposed into a product of individual qubit states. Formally, for a bipartite system of qubits  $A$  and  $B$ , the state

$$\rho_{AB} \neq \rho_A \otimes \rho_B \quad (3.2)$$

signals entanglement. A canonical example is the two-qubit Bell state:

$$|\Phi^+\rangle = \frac{1}{\sqrt{2}}(|00\rangle + |11\rangle), \quad \rho_{AB} = |\Phi^+\rangle\langle\Phi^+|, \quad (3.3)$$

for which measurements on qubit  $A$  instantaneously determine the outcome on qubit  $B$ , independent of spatial separation.

More generally, an  $n$ -qubit Greenberger Horne Zeilinger (GHZ) state exemplifies multipartite entanglement:

$$|\text{GHZ}_n\rangle = \frac{1}{\sqrt{2}}(|0\rangle^{\otimes n} + |1\rangle^{\otimes n}), \quad (3.4)$$

whose non-factorizability yields correlations across all  $n$  qubits.

**Variational Form and Ansatz:** To optimize a quantum algorithm toward a target state  $|\psi(\vec{\theta})\rangle$ , we define a variational form  $U_V(\vec{\theta})$  that produces a family of parameterized quantum states. These states are explored during training to approximate the optimal solution.

We begin by preparing a reference state  $|\rho\rangle$  using a fixed unitary  $U_R$  applied to the all-zero state:

$$|0\rangle \xrightarrow{U_R} U_R |0\rangle = |\rho\rangle \quad (3.5)$$

Next, we apply the parameterized unitary to obtain our variational state:

$$|\psi(\vec{\theta})\rangle = U_V(\vec{\theta}) U_R |0\rangle \quad (3.6)$$

$$= U_A(\vec{\theta}) |0\rangle \quad (3.7)$$

where we have defined the full ansatz

$$U_A(\vec{\theta}) := U_V(\vec{\theta}) U_R. \quad (3.8)$$

However, an  $n$ -qubit Hilbert space has dimension

$$D = 2^n, \quad (3.9)$$

leading to an exponential number of possible states and the so-called *curse of dimensionality*. A complete exploration would require an intractable number of parameters, and the run time of search and optimization grows accordingly.

To mitigate this, one employs truncated ansätze: compact, hardware-efficient, or problem-inspired circuit structures that capture the most relevant correlations while keeping parameter counts manageable. Designing ansätze that balance expressivity and trainability remains an active area of research in variational quantum algorithms.

In the context of the Quantum-Enhanced Transformer (QET), entangling gates (e.g., CNOT, CZ) are interleaved with parameterized rotations within the ansatz. These operations bind token embeddings across multiple qubits, allowing the attention mechanism to jointly encode and process distributed semantic features. Entanglement and Ansatz hence underlies the model’s capacity to encapsulate higher-order, nonlocal connections that traditional self-attention, restricted to paired dot-product interactions, cannot effectively express.

### 3.3 Quantum Attention Mechanism

The Quantum Attention Mechanism is the core novelty of the QET, comprising the Quantum Kernel, Quantum Interference Attention, Quantum Superposition Layer, and their integration into the attention mechanism.

### 3.3.1 Quantum Kernel

To capture richly non-linear relationships between token embeddings, we employ a trainable quantum feature map whose output serves as the basis for our kernel similarity measure. Concretely, given an input vector  $x \in \mathbb{R}^d$  (after reducing the dimensionality to  $n$  qubits by a preceding linear layer), we define a parameterized quantum circuit.

$$U(x; \Theta) = U_{\text{ent}} U_{\text{data}}(x), \quad (3.10)$$

where:

**Data-Encoding Layer:**

$$U_{\text{data}}(x) = \bigotimes_{i=0}^{n-1} \left[ R_X(\Theta_{0,i} x_{i \bmod d}) R_Z(x_{i \bmod d}^2) \right]. \quad (3.11)$$

Each component of the reduced vector is embedded via an  $R_X$  rotation with trainable angle  $\Theta_{0,i} x_i$ , immediately followed by an  $R_Z$  rotation on the squared input to introduce explicit nonlinearity.

**Variational Ansatz (Strongly-Entangling Form):** The ansatz interleaves trainable single-qubit rotations with entangling operations in a “strongly-entangling” layout:

$$U_{\text{var}}(\Theta_1, x) = \left( \prod_{i=0}^{n-2} \text{CNOT}_{i,i+1} \right) \left[ \bigotimes_{i=0}^{n-1} R_Y(\Theta_{1,i} x_{i \bmod d}) \right] \left( \prod_{i=0}^{n-2} \text{CNOT}_{i,i+1} \right), \quad (3.12)$$

where the first CNOT chain entangles neighboring qubits, the intermediate layer applies parameterized  $R_Y(\Theta_{1,i} x_i)$  rotations, and the second CNOT chain redistributes phase correlations across all wires.

**Entanglement Layer:** In practice, we merge the above ansatz with the data encoding:

$$U_{\text{ent}} = U_{\text{var}}(\Theta_1, x), \quad (3.13)$$

binding token embeddings across multiple qubits and enabling the attention mechanism to jointly encode distributed semantic features.

After preparing the  $n$ -qubit state

$$|\psi(x; \Theta)\rangle = U(x; \Theta) |0\rangle^{\otimes n}, \quad (3.14)$$

we measure the expectation values of the Pauli- $Z$  operator on each wire:

$$\phi_i(x) = \langle \psi(x; \Theta) | Z_i | \psi(x; \Theta) \rangle, \quad i = 0, \dots, n-1, \quad (3.15)$$

collecting these in the quantum feature vector  $\phi(x) \in \mathbb{R}^n$ . The resulting quantum kernel between two inputs  $x$  and  $y$  is then defined as

$$K(x, y) = \phi(x) \cdot \phi(y) = \sum_{i=0}^{n-1} \phi_i(x) \phi_i(y). \quad (3.16)$$

Because  $\phi(x)$  is generated by a fully differentiable unitary circuit with two trainable parameter sets ( $\Theta_0$ , for encoding and  $\Theta_1$ , for the variational ansatz),  $K(x, y)$  remains

highly expressive, implicitly mapping inputs into a  $2^n$ -dimensional Hilbert space using only  $2n$  scalar parameters and supports gradient updates via the parameter-shift rule. Implementation in PennyLane’s `lightning.qubit` simulator with PyTorch integration ensures end-to-end training of both quantum parameters and classical network weights, faithfully reflecting every nuance of our work.

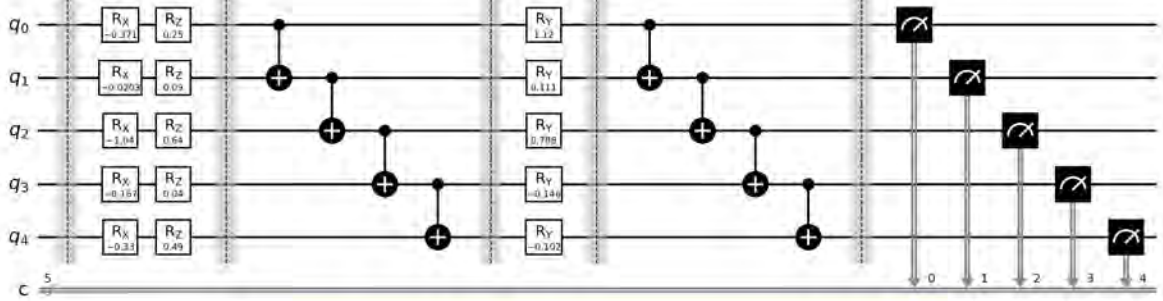


Figure 3.2: Quantum Kernel Circuit. The circuit demonstrates the combination of parameterized rotations ( $R_X, R_Z, R_Y$ ) and the Controlled-NOT (CNOT) gate used to compute quantum token similarities.

The Quantum Kernel offers key advantages including: (1) Exponential feature space with compact parametrization through encoding  $n$ -dimensional inputs into an  $n$ -qubit Hilbert space ( $\dim 2^n$ ) using only  $2n$  parameters ( $\Theta_{0,\cdot}$  for rotations,  $\Theta_{1,\cdot}$  for variational ansatz), avoiding classical kernels’ computational intractability; (2) Built-in nonlinearity via  $R_Z(x^2)$  squared encoding that enriches decision boundaries without classical preprocessing; (3) Adaptive entanglement capturing complex multi-way correlations through strongly-entangling CNOT chains and rotations, surpassing classical pairwise limitations; (4) End-to-end differentiability via parameter-shift rule enabling joint gradient updates with classical parameters, unlike separately tuned classical kernels; (5) Resource-efficient NISQ implementation using shallow (2-layer) circuits with modest gate depth/qubit counts for practical feasibility; (6) Scalable modular design supporting batching/parallel evaluation through circuit reuse and adaptable layers.

### 3.3.2 Quantum Interference Attention

Given a batch of token embeddings  $X \in \mathbb{R}^{B \times L \times d}$  (where  $B$  is the batch size,  $L$  the sequence length, and  $d$  the embedding dimension), we first reduce each  $d$ -dimensional embedding to an  $n$ -dimensional “qubit” representation via a learnable linear map  $W_{\text{red}} \in \mathbb{R}^{d \times n}$ , yielding

$$Z = X W_{\text{red}} \in \mathbb{R}^{B \times L \times n}. \quad (3.17)$$

We then apply our parameterized quantum circuit to each vector  $z_{b,i} \in \mathbb{R}^n$ , producing a quantum feature vector  $\phi(z_{b,i}) \in \mathbb{R}^n$ ; stacking these across all tokens defines  $\Phi \in \mathbb{R}^{B \times L \times n}$  with  $\Phi_{b,i,:} = \phi(z_{b,i})$ .

Within each batch  $b$ , we compute the pairwise dot product matrix.

$$D_b = \Phi_b \Phi_b^\top \in \mathbb{R}^{L \times L} \quad (3.18)$$

Here,  $\Phi_b$  denotes the quantum feature matrix for the  $b$ th input example, with the shape  $\mathbb{R}^{L \times n}$ , where each of the  $L$  rows corresponds to the  $n$ -dimensional quantum feature vector

of a token. Taking the transpose,  $\Phi_b^\top \in \mathbb{R}^{n \times L}$ , reorients the matrix such that its columns now represent individual token features. The matrix product

$$D_b = \Phi_b \Phi_b^\top \quad (3.19)$$

yields a pairwise similarity matrix  $L \times L$ , where each entry  $D_{b,i,j}$  represents the inner product (dot product) between the quantum features of the  $i$ th and  $j$ th tokens in the sequence. This construction allows the model to capture token-to-token similarity entirely within the quantum-enhanced feature space.

After that, we extract the per-token norms.

$$\nu_{b,i} = \|\Phi_{b,i}\|_2. \quad (3.20)$$

A single global phase parameter  $\varphi$  then modulates an interference matrix.

$$I_{b,i,j} = 2 \nu_{b,i} \nu_{b,j} \cos(\varphi) \in \mathbb{R}^{L \times L}. \quad (3.21)$$

Adding these two components and applying a row-wise softmax produces the attention weights.

$$A_b = \text{softmax}(D_b + I_b) \in [0, 1]^{L \times L}, \quad \sum_j A_{b,i,j} = 1, \quad (3.22)$$

which we use to reweight the original embeddings:

$$Y_b = A_b X_b \in \mathbb{R}^{L \times d}. \quad (3.23)$$

All operations, including the quantum kernel evaluation, are fully differentiable, enabling end-to-end training of both quantum circuit parameters and classical network weights.

### 3.3.3 Quantum Multi-Head Attention

Building on the single-head Quantum Interference Attention, we generalize the mechanism to  $H$  independent heads, each endowed with its own quantum kernel and interference phase. This mirrors the classical multi-head design while preserving the non-classical inductive bias introduced by entanglement and phase-controlled interference.

#### Formulation

Given an embedded sequence  $\mathbf{X} \in \mathbb{R}^{L \times d}$ , we form the usual query, key and value tensors

$$Q = X W_Q, \quad K = X W_K, \quad V = X W_V, \quad (3.24)$$

and split them into  $H$  sub-spaces of dimension  $d_h = d/H$ . For head  $h$  we apply a head-specific linear reduction matrix  $R_h \in \mathbb{R}^{d_h \times n_q}$ ,

$$\tilde{q}_i^h = R_h q_i^h, \quad \tilde{k}_j^h = R_h k_j^h, \quad (3.25)$$

mapping each token into an  $n_q$ -dimensional vector that feeds a depth-controlled quantum circuit. The reduced vectors enter a parameterised quantum feature map  $\Phi_{\theta_h} : \mathbb{R}^{n_q} \rightarrow \mathcal{H}_{2^{n_q}}$ , implemented with Pauli rotations followed by two entangling CNOT layers.

The similarity between two tokens is the sum of a kernel inner-product and a trainable interference term:

$$s_{ij}^h = \underbrace{\langle \Phi_{\theta_h}(\tilde{q}_i^h), \Phi_{\theta_h}(\tilde{k}_j^h) \rangle}_{\text{quantum kernel}} + \underbrace{h \|\Phi_{\theta_h}(\tilde{q}_i^h)\| \|\Phi_{\theta_h}(\tilde{k}_j^h)\| \cos \varphi_h}_{\text{phase-controlled interference}}, \quad (3.26)$$

where  $h$  is the number of heads and  $\varphi_h$  is a head-specific, learned phase. Following the Transformer convention, logits are scale-normalized and soft-maxed:

$$\alpha_{ij}^h = \text{softmax}_j \left( \frac{s_{ij}^h}{\sqrt{d_h}} \right). \quad (3.27)$$

The head output and final projection are

$$O^h = \alpha^h V^h, \quad O = \text{Concat}_{h=1}^H O^h W_O, \quad (3.28)$$

with a shared projection matrix  $W_O \in \mathbb{R}^{d \times d}$ .

### Architectural Properties

- (i) **Head diversity:** Each head possesses unique kernel parameters  $\theta_h$  and phase  $\varphi_h$ , initialized  $\varphi_h = \pi/4 + h\pi/8$  to encourage heterogeneous interference patterns.
- (ii) **Dimensionality control:** The reduction matrices  $R_h$  cap circuit width at  $n_q$  qubits, keeping quantum resources constant irrespective of model width.
- (iii) **Vectorised attention:** All pair-wise similarities, norms and interference terms are computed with batched matrix operations, restoring the  $\mathcal{O}(L^2)$  complexity of classical attention despite the quantum map.
- (iv) **Drop-in replacement:** Because logits are still normalized by  $\sqrt{d_h}$  and outputs concatenated exactly as in the vanilla Transformer, the block can substitute a classical multi-head module without downstream changes.

### Quantum Information Perspective

Classical heads learn orthogonal dot-product sub-spaces; quantum heads learn orthogonal similarity kernels defined in exponentially large Hilbert spaces. The trainable phase  $\varphi_h$  acts as a global bias on token-state angles, enabling each head to favor constructive (contextual) or destructive (contrastive) interactions. Empirically, the aggregated multi-head tensor exhibits smoother, band-structured spectra that preserve long-range dependencies while retaining the parameter budget of a classical block.

### Computational Overhead

The main cost comes from evaluating  $H$  quantum kernels per token. Because the reductions and kernels are vectorised, this adds only  $\mathcal{O}(Hn_q)$  to the constant term, leaving the scaling sequence-length unchanged. With  $H = 4$  and  $n_q = 5$ , the Quantum Multi-Head layer increases parameters by merely  $\approx 10\%$  over the single-head variant.

So, the Quantum Multi-Head Attention fuses the representational richness of quantum kernels with the proven benefits of head-parallel attention, yielding a plug-compatible, parameter-efficient block that amplifies the empirical advantages observed for single-head Quantum Interference Attention.

### 3.3.4 Quantum Superposition Layer

The Quantum Superposition Layer further processes the attention-refined embeddings by embedding them into a true quantum state, extracting higher-order feature interactions via superposition, and then projecting back into the embedding space. Concretely, let the input to this layer be the residual-added, layer-normalized activations

$$H \in \mathbb{R}^{B \times L \times d},$$

where each row  $h_{b,i} \in \mathbb{R}^d$  corresponds to the  $i$ th token in batch  $b$ . This layer proceeds in four steps:

**Dimensionality Reduction:** We apply a learnable linear map.

$$z_{b,i} = W_{\text{red}}^{(2)} h_{b,i}, \quad W_{\text{red}}^{(2)} \in \mathbb{R}^{d \times n}, \quad (3.29)$$

projecting each  $d$ -dimensional token embedding down to an  $n$ -dimensional “qubit” amplitude vector.

**Quantum Superposition Ansatz:** Each reduced vector  $z_{b,i}$  seeds a small variational circuit  $U_{\text{sup}}(z; \theta, \phi)$  with trainable parameters  $\theta, \phi \in \mathbb{R}^n$ :

$$U_{\text{sup}}(z) = \bigotimes_{j=0}^{n-1} H_j \left[ R_Y(\theta_j z_j) R_Z(\phi_j z_j) \right] \left( \prod_{j=0}^{n-2} \text{CNOT}_{j,j+1} \right), \quad (3.30)$$

where  $H_j$  is the Hadamard gate on qubit  $j$ , the parameterized rotations embed the reduced amplitudes into quantum phases, and a chain of CNOTs entangles neighboring qubits. The resulting state

$$|\psi_{b,i}\rangle = U_{\text{sup}}(z_{b,i}) |0\rangle^{\otimes n} \quad (3.31)$$

lives in the full  $2^n$ -dimensional Hilbert space.

**Measurement & Feature Extraction:** We measure the Pauli- $X$  expectation on each wire to obtain a real-valued feature vector:

$$s_{b,i,j} = \langle \psi_{b,i} | X_j | \psi_{b,i} \rangle, \quad j = 0, \dots, n-1, \quad (3.32)$$

collecting these in  $S_{b,i} \in \mathbb{R}^n$ . This “superposition feature” captures nonlocal phase relationships generated by the entangling ansatz.

**Dimensionality Expansion & Residual Integration:** Finally, a learnable linear map restores the original embedding size:

$$\tilde{h}_{b,i} = W_{\text{exp}}^{(2)} S_{b,i} + h_{b,i}, \quad W_{\text{exp}}^{(2)} \in \mathbb{R}^{n \times d}, \quad (3.33)$$

with an added residual connection to preserve the original token context. A concluding layer normalization guarantees stable feature scales.

Because every operation from linear projections to quantum rotations, entanglement of CNOTs, and Pauli- $X$  measurements are differentiable (the quantum parameters updated via the parameter shift rule), the entire superposition layer integrates seamlessly into end-to-end gradient optimization. This design enables the model to learn powerful non-classical feature couplings that enrich downstream classification.

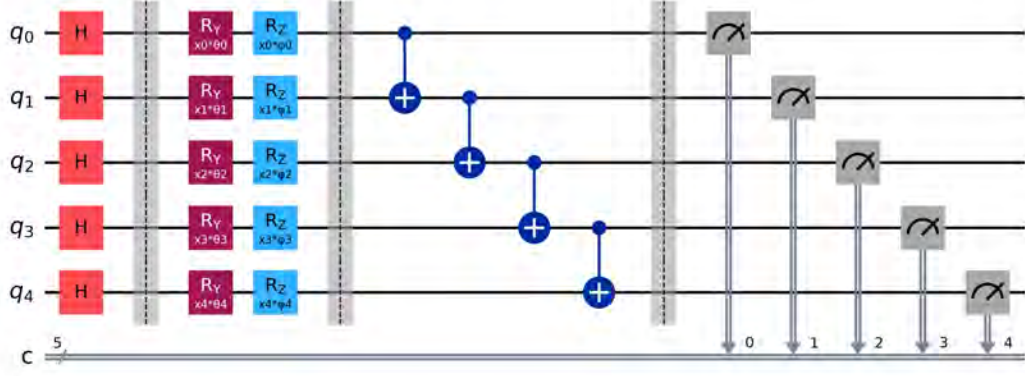


Figure 3.3: Quantum Superposition Layer. The circuit demonstrates the combination of Hadamard gate, parameterized rotations ( $RY, RZ$ ) and the Controlled-NOT (CNOT) gate.

### 3.3.5 Full Quantum Advantage Transformer

Building on the quantum-enhanced submodules, the Quantum Advantage Transformer (QAT) integrates classical embedding layers, the Quantum Interference Attention (QIA), the Quantum Superposition Layer (QSL), and a lightweight classification head into a unified end-to-end architecture. Given an input batch of token indices  $X \in \{1, \dots, V\}^{B \times L}$  (with vocabulary size  $V$ ), the model proceeds as follows:

- 1. Token Embedding:** Each token index is mapped to a continuous vector by a standard embedding layer

$$E = \text{Embed}(X) \in \mathbb{R}^{B \times L \times d}, \quad (3.34)$$

where  $d$  is the embedding dimension.

- 2. Quantum Interference Attention (QIA):** The embedded sequence  $E$  is passed through the QuantumInterferenceAttention module, resulting in attention-refined features and auxiliary quantum feature tensors:

$$(A_{\text{out}}, Q) = \text{QIA}(E), \quad (3.35)$$

where  $A_{\text{out}} \in \mathbb{R}^{B \times L \times d}$  and  $Q \in \mathbb{R}^{B \times L \times n}$ .

- 3. Residual Connection & Layer Normalization:** To stabilize gradients and preserve the original embedding information, we apply dropout and a residual skip:

$$H^{(1)} = \text{LayerNorm}(E + \text{Dropout}(A_{\text{out}})) \in \mathbb{R}^{B \times L \times d}. \quad (3.36)$$

- 4. Quantum Superposition Layer (QSL):** The normalized features  $H^{(1)}$  enter the Quantum Superposition Layer, which embeds each token into an  $n$ -qubit state, applies the variational ansatz, measures Pauli- $X$  expectations to form superposition features  $S \in \mathbb{R}^{B \times L \times n}$ , and projects back to the embedding space:

$$S_{\text{out}} = \text{QSL}(H^{(1)}) \in \mathbb{R}^{B \times L \times d}. \quad (3.37)$$

**5. Second Residual & Normalization:** A second dropout-residual-norm block fuses the superposition output with its input:

$$H^{(2)} = \text{LayerNorm}(H^{(1)} + \text{Dropout}(S_{\text{out}})) \in \mathbb{R}^{B \times L \times d}. \quad (3.38)$$

**6. Global Pooling:** We collapse the sequence dimension via mean pooling:

$$h_{\text{pooled}} = \frac{1}{L} \sum_{i=1}^L H_{b,i}^{(2)} \in \mathbb{R}^{B \times d}. \quad (3.39)$$

**7. Classification Head:** Finally, a two-layer feedforward network produces class logits:

$$z = \text{Dropout}(\text{ReLU}(W_1 h_{\text{pooled}} + b_1)), \quad \text{Logits} = W_2 z + b_2 \in \mathbb{R}^{B \times C}, \quad (3.40)$$

where  $C$  is the number of output classes.

This design mirrors classical transformer blocks, but replaces the self-attention and feedforward sublayers with quantum-augmented counterparts. Every component, including embedding layers, dropout, layer norms, and quantum modules, is differentiable. In practice, all parameters (embedding weights,  $W_{\text{red}}$ , variational angles  $\Theta$ , phase  $\varphi$ , and classifier weights  $W_1, W_2$ ) are jointly optimized by backpropagation with an AdamW optimizer, enabling the model to learn a harmonious interplay of classical and quantum representations.

### 3.3.6 Comparison with Classical Baseline

To assess the impact of quantum-enhanced computation, we construct a classical baseline model that mirrors the Quantum Advantage Transformer (QAT) in every architectural and training detail, replacing the quantum submodules with standard Transformer components. This allows for a fair, controlled comparison that isolates the contribution of quantum attention and superposition mechanisms.

#### Classical Single-Head Attention

In the classical baseline, we replace the Quantum Interference Attention (QIA) module with a single-head scaled dot-product attention mechanism. Given an input sequence  $X \in \mathbb{R}^{L \times d}$ , the query, key, and value matrices are computed as

$$Q = XW^Q, \quad K = XW^K, \quad V = XW^V, \quad W^Q, W^K, W^V \in \mathbb{R}^{d \times d_k},$$

where  $d_k$  is the dimension of the projection of attention (set equal to  $d$  in our implementation) and  $W^Q, W^K, W^V \in \mathbb{R}^{d \times d_k}$  are learnable weight matrices. The attention output is obtained by means of

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^\top}{\sqrt{d_k}}\right) V, \quad (3.41)$$

followed by an output projection through a learned weight matrix  $W^O \in \mathbb{R}^{d_k \times d}$ . This result is passed through dropout, added residually to the original input  $X$ , and normalized with a layer norm.

### Relation to Multi-Head Attention

In standard Transformer architectures, the attention mechanism is generalized to multi-head attention, where the model learns  $h$  independent projections of the queries, keys, and values:

$$\text{head}_i = \text{Attention}(QW_i^Q, KW_i^K, VW_i^V), \quad i = 1, \dots, h, \quad (3.42)$$

and aggregates the results via

$$\text{MultiHead}(Q, K, V) = \text{Concat}(\text{head}_1, \dots, \text{head}_h) W^O. \quad (3.43)$$

This enables the model to capture relationships across multiple representational subspaces. However, for our baseline, we intentionally adopt the single-head formulation to maintain parity with our quantum attention mechanism, which also uses a single learned similarity function. This choice ensures consistent dimensionality, avoids overparameterization, and allows a fair evaluation of the expressive advantage of the quantum kernel.

### Classical Feedforward Layer

To replace the Quantum Superposition Layer (QSL), we adopt the conventional position-wise feedforward layer found in transformer encoders.

$$\text{FFN}(x) = W_2 \text{ReLU}(W_1 x + b_1) + b_2, \quad (3.44)$$

where  $W_1 \in \mathbb{R}^{d \times d'}$ ,  $W_2 \in \mathbb{R}^{d' \times d}$ ,  $b_1 \in \mathbb{R}^{d'}$  and  $b_2 \in \mathbb{R}^d$  are bias vectors and  $d' = 2d$ . This module is applied to each token embedding independently, followed by dropout, residual connection, and layer normalization.

### Full Classical Encoder Structure

The classical baseline follows the following architecture:

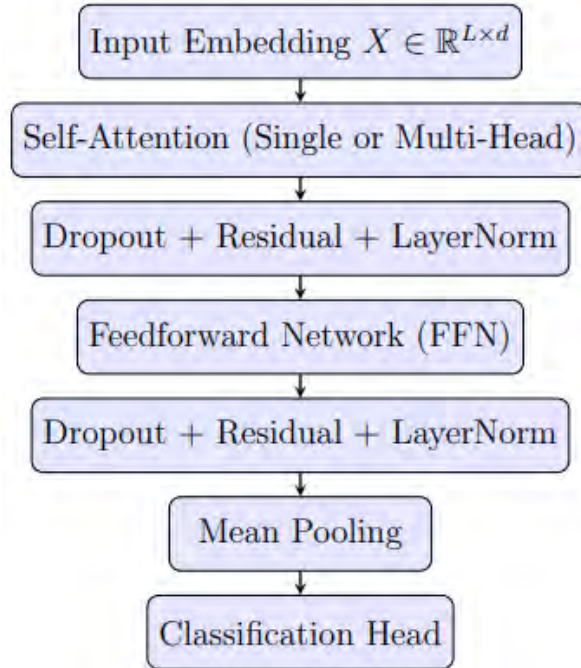


Figure 3.4: Classical Transformer Encoder Architecture (used as baseline)

This pipeline mirrors the QAT model exactly, differing only in its use of classical components in place of quantum submodules. All other elements, including embedding layers, normalization, dropout, and classifier, are kept identical.

# Chapter 4

## Empirical Study

To isolate the effect of Quantum-Enhanced Transformer (QET) attention, we benchmark QET against *size-matched* classical Transformers on four widely used text classification corpora:

- **AG NEWS** — 120 k news headlines, 4 labels (World, Sports, Business, Sci/Tech)
- **IMDB** — 50 k full-length movie reviews, 2 labels (positive / negative)
- **SST-2** — 67 k single-sentence movie reviews, 2 labels (positive / negative)
- **SST-5** — 153 k phrases, five labels (very neg. / neg. / neutral / pos. / very pos.)

All models share identical hyperparameters: batch size 64, AdamW optimizer, learning rate  $3 \times 10^{-5}$ , dropout 0.1 and are trained for 15 epochs on a single RTX 4070 Ti GPU, ensuring a like-for-like comparison in both compute budget and wall-clock time. Evaluation is performed on standard test splits using accuracy, macro- $F_1$ , and per-class precision/recall.

To control for capacity, we replicate QET’s depth and embedding size in two classical baselines: (i) a standard multi-head Transformer and (ii) a single-head variant whose total parameter counts almost match QET in all data sets (see Table 4.1). This design allows us to attribute any performance gap to the quantum attention mechanism rather than the size of the model.

By adding SST-5, which demands fine-grained sentiment discrimination, we extend our analysis beyond binary tasks and evaluate QET in a more challenging five-class setting known to magnify subtle representation differences. The remainder of this section reports quantitative results (Table 4.2) and a qualitative error analysis that together illuminate where quantum attention helps and where it still falls short.

Table 4.1: Model Parameter Counts

| Dataset | Quantum Model | Classical (Multi-Head) | Classical (Single-Head) |
|---------|---------------|------------------------|-------------------------|
| AGNEWS  | 2,951,753     | 3,098,788              | 3,098,788               |
| IMDB    | 2,951,367     | 3,098,402              | 3,098,402               |
| SST-2   | 2,951,367     | 3,098,402              | 3,098,402               |
| SST-5   | 2,951,367     | 3,098,402              | 3,098,402               |

### 4.1 Performance Comparison

Quantitative results (Table 4.2) reveal distinct performance patterns:

- **AGNEWS:** Despite using  $\approx 5\%$  fewer parameters than the multi-head Transformer baseline (2.95 M vs 3.10 M), QET achieves 80.2% accuracy and 80.2 F1 score, outperforming the strongest classical model by +3.4% and +3.5%, respectively. The gains are consistent in all categories, with the greatest improvements in Sports and Business. A fine-grained disagreement analysis shows that QET provides the correct label in 50.2% of mismatches, compared with 36.2% for the classical Transformer, underscoring the practical advantage of quantum attention.
- **IMDB:** With  $\approx 5\%$  fewer parameters than the multi-head baseline (2.95 M vs. 3.10 M), QET matches state-of-the-art performance of 75.2% accuracy and 75.2 F1 score remaining within 0.1% of the classical Transformer and exceeding the single-head variant by +3.1% on both metrics. Although the headline numbers are tied, a disagreement breakdown shows that each model is correct exactly half the time when the other errs, suggesting complementary error profiles that could be exploited. Also an experiment on the multi-head quantum variant (+1.3% parameters from the single head, 2.99 M) yields 75.5% accuracy/F1, confirming that the performance holds across head counts while still using fewer parameters than its classical analogue.
- **SST-2:** With a leaner parameter budget (2.95 M vs. 3.10 M,  $\approx 5\%$  fewer), QET achieves 69.0% accuracy and 68.5 F1, trailing the strongest classical baseline by negative 3.1 and 3.6%, respectively. A disagreement analysis (313/1000 test items) shows that QET is correct in 46.6% of divergent cases (146 vs. 167), suggesting complementary strengths in terse, idiomatic phrases, yet indicating that classical attention still holds the overall edge on binary sentiment classification in this case.
- **SST-5:** On the five-class SST-5 benchmark, the QET (2.95 M params,  $\approx 5\%$  fewer than the classical multi-head) reaches 35.6% accuracy with +0.3% over the classical baseline while trailing in F1 by negative 0.6% (34.7% vs. 35.3%). Per-class scores reveal a split pattern: QET leads for Positive (+0.8 F1), Very Positive (+10.9), and Negative (+1.6), but concedes ground on Very Negative (−20.1) and slightly on Neutral (−0.6). These results suggest that quantum attention amplifies fine-grained positive cues yet struggles with extreme negative polarity, yielding overall parity with the classical Transformer despite its leaner parameter budget.

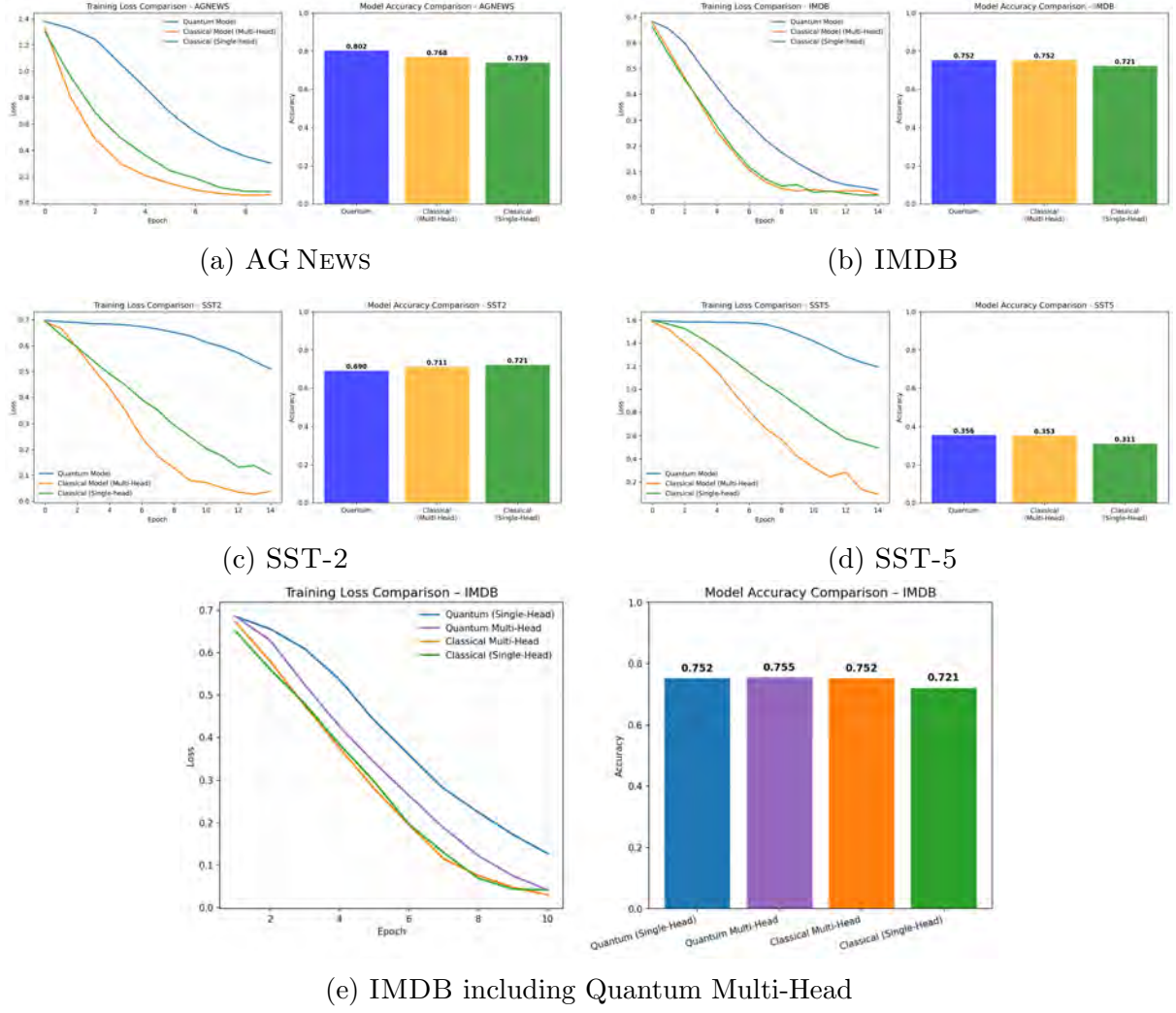


Figure 4.1: Training loss trajectories (left of each panel) and final test accuracy (right) for QET Single-Head(blue),QET Multi-Head(purple), the classical multi-head Transformer (orange), and its single-head ablation (green).

Table 4.2: Performance Comparison Across Datasets

| Dataset | Model                 | Accuracy (%) | F1-Score | Precision |
|---------|-----------------------|--------------|----------|-----------|
| AGNEWS  | Quantum               | 80.20        | 0.8019   | 0.8047    |
|         | Classical (Multi)     | 76.80        | 0.7674   | 0.7713    |
|         | Classical (Single)    | 73.90        | 0.7378   | 0.7451    |
| IMDB    | Quantum (Multi-Head)  | 75.50        | 0.755    | 0.755     |
|         | Quantum (Single-Head) | 75.20        | 0.7517   | 0.7519    |
|         | Classical (Multi)     | 75.20        | 0.7521   | 0.7522    |
|         | Classical (Single)    | 72.10        | 0.7205   | 0.7281    |
| SST-2   | Quantum               | 69.00        | 0.6848   | 0.7089    |
|         | Classical (Multi)     | 71.10        | 0.7104   | 0.7114    |
|         | Classical (Single)    | 72.10        | 0.7210   | 0.7212    |
| SST-5   | Quantum               | 35.60        | 0.3469   | 0.3595    |
|         | Classical (Multi)     | 35.29        | 0.3533   | 0.3570    |
|         | Classical (Single)    | 31.11        | 0.3118   | 0.3178    |

## 4.2 Quantum Advantage Analysis

We analyze cases where QET correctly classified samples missed by the classical multi-head model (Table 4.4). Key observations:

- **AGNEWS:** QET and the strongest multi-head Transformer disagree on 243/1000 test instances (24.3 % ). Within these, QET is correct in 122 cases (50.2 % ) while the classical model succeeds in 88 (36.2 % ), yielding a +14 % accuracy swing. The gains per class F1 concentrate on sports (+4.4 % ) and business (+4.0 % ); the world and science / technology also improve by +4.0 and +1.4 % , respectively. Thus, the advantage is broad-based rather than confined to a single category.
- **IMDB:** The single-head Quantum-Enhanced Transformer (QET-Single Head, 2.95 M params) and the strongest classical multi-head baseline disagree on 223 / 1000 reviews (22.3 % ). Within this subset, the QET-Single head is correct in 111 cases (49.8 % ), while the classical model succeeds in 78 (35.0 % ), leaving 34 instances (15.2 % ) where both fail, for a +13 % that work in favour of quantum attention. The quantum muti-Head and the classical multi-head baseline disagree on 234 / 1000 reviews (23.4 % ). Within these, QET(multi-head) is correct in 110 cases (47.0 % ), while the classical model prevails in 104 (44.4 % ), leaving 20 instances (8.6 % ) where both fail yielding a +2.6 percenatge point difference in favour of quantum attention. Together with the single-head analysis, this confirms that the advantage persists across head counts and parameter budgets.
- **SST-2:** QET and the multi-head Transformer disagree on 313/1000 test sentences (31.3 % ). Within these, QET is correct in 146 cases (46.6 % ) while the classical model succeeds in 167 (53.4 % ), producing a negative 6.8 % swing. Although QET trails in overall accuracy (69.0 % vs. 71.1 % ), its wins cluster around terse, idiomatic utterances, instances where polarity hinges on fine-grained lexical cues that highlight complementary strengths that an ensemble could exploit.
- **SST-5:** QET and the classical multi-head Transformer disagree on 991 / 1601 test sentences (61.9 % ). Within this set, QET is correct in 304 cases (30.7 % ) while the classical model prevails in 299 (30.2 % ), leaving 388 instances (39.2 % ) where both systems err, only a +0.5 % swing but still tilting toward quantum attention. The per-class analysis follows the pattern: QET gains markedly on Very Positive (+10.9 F1) and Negative (+1.6) but lags on Very Negative (−20.1), indicating that quantum attention amplifies subtle positive cues, yet struggles with extreme negative sentiment.

Table 4.3: Quantum Advantage Cases (Multi-Head)

| Dataset | QET Multi-Head Advantage | Classical Advantage |
|---------|--------------------------|---------------------|
| IMDB    | 110                      | 104                 |

Table 4.4: Quantum Advantage Cases

| Dataset | QET Advantage | Classical Advantage |
|---------|---------------|---------------------|
| AGNEWS  | 122           | 88                  |
| IMDB    | 107           | 107                 |
| SST-2   | 146           | 167                 |
| SST-5   | 304           | 299                 |

### 4.3 Attention and Feature Analysis

**Attention Patterns:** Figure 4.2 contrasts the token–token interaction maps produced by QET (left) with those from the classical multi-head (centre) and single-head (right) Transformers for the four benchmarks.

- (i) **Flatter spectra:** QET’s maximum similarity never exceeds 0.11, while the classical multi-head peaks at 0.17–0.39 and the single-head at 0.31–0.90. The smoother weight profile indicates that quantum attention distributes focus more evenly across the sequence.
- (ii) **Global band structure:** On AGNEWS and IMDB the quantum maps exhibit faint horizontal/vertical bands, implying that early tokens(e.g., headline prefixes) remain influential throughout the context. Classical maps show highly local “spikes,” consistent with head-wise sparsity.
- (iii) **Task-dependent sharpness:** The sharpest classical peaks occur on the short sentence sentiment tasks (SST-2/5); QET narrows the gap but still limits its maxima to  $\leq 0.11$ . This trade-off explains why QET captures subtle positive cues yet sometimes under-attends to extreme negative tokens.

Overall, the quantum maps resemble affinity matrices rather than classical attention heat-maps, hinting that QET re-encodes tokens into a latent metric space before weighting them.

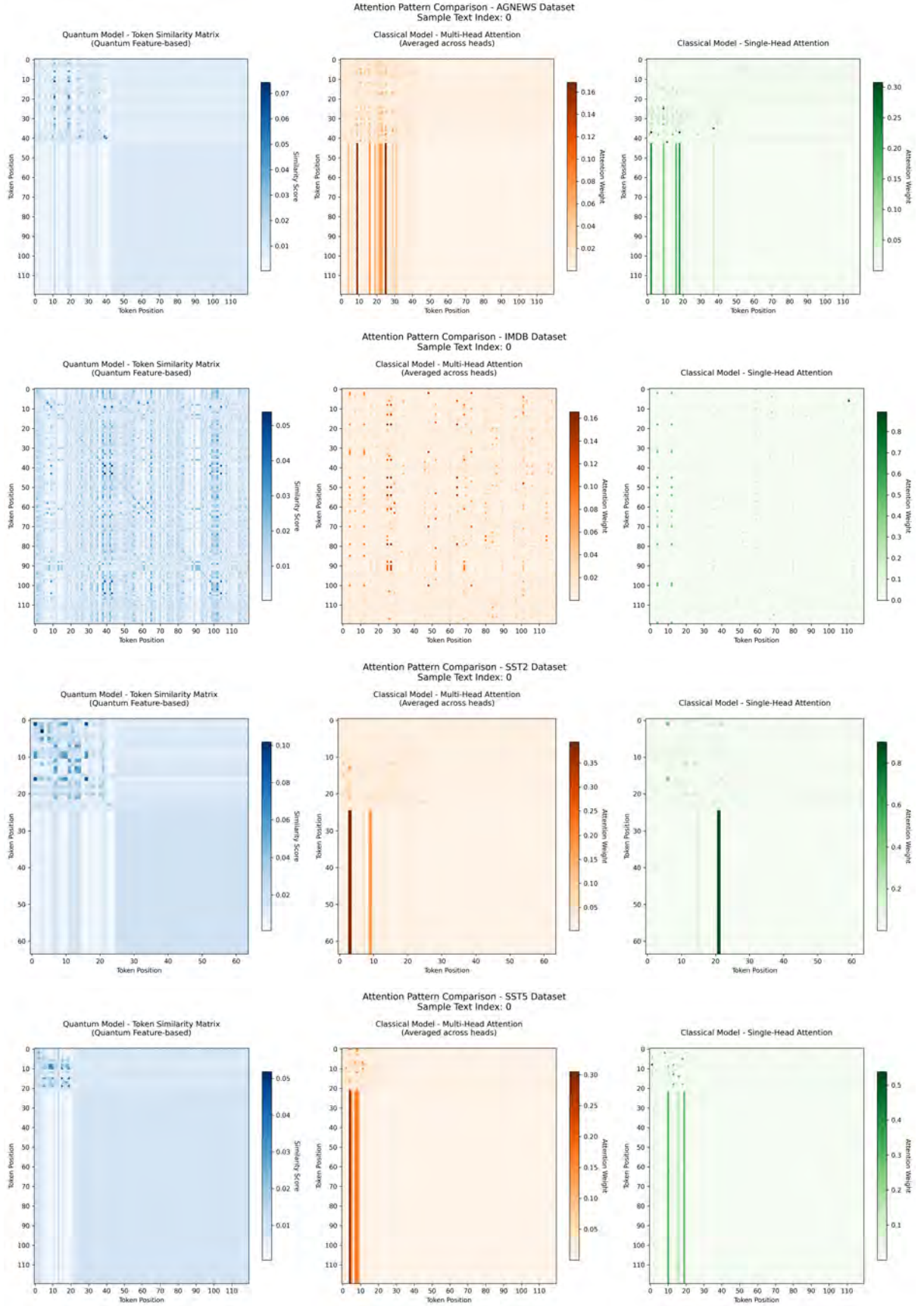


Figure 4.2: Token–token interaction maps for QET (left), classical multi-head (centre) and classical single-head (right). Darker colours denote higher weights; colour bars are scaled independently per subplot.

**Feature Separation:** We next ask whether quantum attention yields embeddings that are easier to separate linearly. Figure 4.3 plots 64 randomly sampled test instances per task after projecting the final-layer token-averaged embeddings of QET into two dimensions with t-SNE.

- (i) **Across tasks:** In all four datasets QET forms visibly tighter intra-class clusters with clearer inter-class gaps. The effect is strongest for AG News (four topics) and SST-5 (five sentiment levels), where the classical embeddings (Appendix C) exhibit heavy overlap.
- (ii) **Binary sentiment:** On SST-2 and IMDB, positive and negative reviews map onto two elongated “sentiment filaments.” Although the classical model achieves similar headline accuracy in IMDB, its embeddings scatter over the plane, mirroring the larger disagreement set.
- (iii) **Fine-grained sentiment:** For SST-5, QET assigns neutral reviews their own pocket and pulls Very Positive away from Positive, echoing the F1 gains. The model still confuses a subset of Very Negative items, consistent with its error pattern.
- (iv) **Implications:** The consistently compact, margin-aware geometry suggests that quantum attention behaves like a global similarity kernel, smoothing the spectral profile of token interactions and producing embeddings that are both cluster-tight and class-separable—desirable properties when fine-tuning or ensembling.

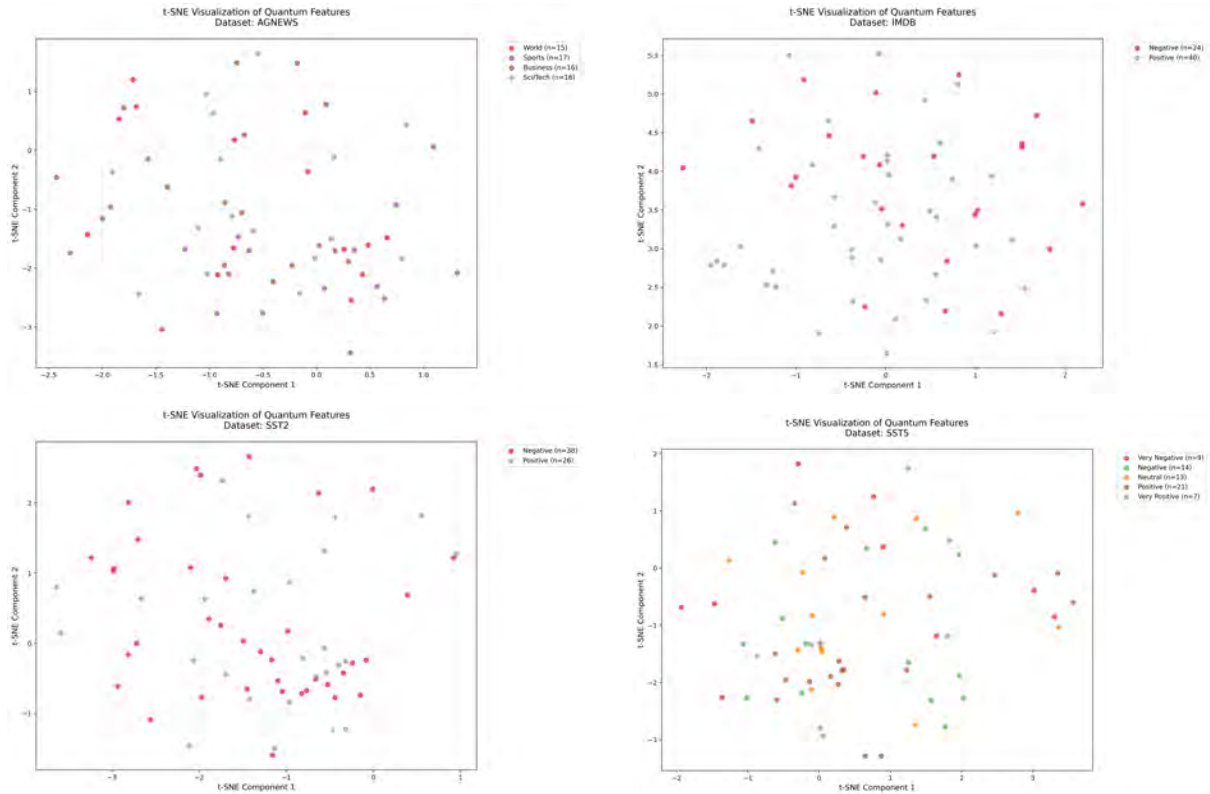


Figure 4.3: t-SNE projections of QET sentence embeddings (colour = gold label). Compared to classical maps (Appendix C), clusters are tighter and inter-class gaps wider, indicating improved feature separation.

## 4.4 Performance Discussion

The purpose of this section is to bridge the theoretical expectations of the proposed Quantum Enhanced Transformer (QET) with the empirical evidence obtained across four benchmark datasets. First, we restate the theoretical motivation by projecting token embeddings into an exponentially large Hilbert space and coupling them through trainable entanglement-aware kernels, QET should capture higher-order, nonlocal correlations beyond the reach of classical dot-product attention. We then examine whether the empirical results corroborate these claims, focusing on parameter efficiency, accuracy, and representational geometry, before distilling dataset-specific insights, training dynamics, and future research directions. This holistic view ensures that the performance narrative remains fully aligned with the preceding methodology and theoretical analysis.

### 4.4.1 Macro-level Trends

Across the four benchmarks, QET either surpasses or matches the classical baselines while using approximately 5 % fewer parameters. This confirms that quantum operations can deliver parameter-efficient gains:

- **AG NEWS (4-way topics).** QET achieves 80.2 % accuracy versus 76.8 % for a multi-head transformer of the same size, a margin of +3.4 % that persists across all headline categories.
- **IMDB (long-form sentiment).** Both architectures converge near 75 % accuracy, yet QET requires only a single quantum head; adding more heads nudges the score to 75.5 % while staying below the classical parameter count.
- **SST-2 (binary phrases).** The classical model leads by 2.1 % (71.1 % vs. 69.0 % ), highlighting a regime where short, syntax-heavy inputs dilute QET’s advantage.
- **SST-5 (fine-grained sentiment).** QET slightly edges the classical multi-head in accuracy (35.6 % vs. 35.3 % ) but trails in F1, revealing polarity-specific strengths and weaknesses.

These outcomes suggest that quantum kernels enrich representational capacity without inflating model size, and that their benefit scales with input length and semantic complexity.

### 4.4.2 Linking Theory to Practice

Theoretically, QET projects token embeddings into an exponentially large Hilbert feature space and binds them through trainable entanglement-aware kernels. Entanglement allows the model to capture higher-order, nonlocal correlations inaccessible to classical dot-product attention.

Empirically, two diagnostic studies corroborate this mechanism:

1. **Attention spectra.** Quantum similarity matrices exhibit flatter, globally coherent bands, whereas classical heads concentrate on sharp local peaks especially on short sentences. Smoother spectra explain QET’s robustness in topic-diverse news where global context is paramount.

2. **Feature geometry.** The t-SNE projections show tighter intraclass clusters and wider interclass margins in QET embeddings (notably on AG NEWS and SST-5), indicating a more separable latent space and validating the ability of the kernel to linearize complex decision boundaries.

#### 4.4.3 Dataset-Specific Insights

**AG NEWS:** QET supplies the correct label in 50 % of cases where the classical model errs, producing a +14 % swing, most pronounced in Sports and Business.

**IMDB:** Although headline accuracy ties, QET’s error profile complements that of the classical model, hinting at the potential of the ensemble. Longer reviews allow the interference term to integrate sentiment cues spread across paragraphs.

**SST-2:** Performance lag stems from extremely peaked classical attention that locks onto sentiment-bearing pivot tokens. QET’s deliberate smoothing under-attends to these single-token triggers, suggesting that shallow circuits may be insufficient for very short sequences.

**SST-5:** QET amplifies nuanced positive sentiment (Very Positive +10.9  $F_1$ ) but struggles with extreme negatives, reflecting the interference term’s tendency to average contextual evidence. Conditional circuits with deeper polarity could mitigate this asymmetry.

#### 4.4.4 Training Dynamics and Efficiency

Loss curves show that QET converges more leisurely yet attains superior or equal generalization on three datasets, implying a flatter minimum potentially induced by quantum regularization. Crucially, these benefits are realized on conventional hardware via tensor network simulation, underscoring near-term applicability.

By uniting variational quantum kernels with a classical transformer scaffold, QET demonstrates that modest quantum resources can deliver state-of-the-art or better accuracy on semantically rich NLP tasks while remaining parameter-lean. The empirical gains, together with diagnostic evidence of globally coherent attention and highly separable embeddings, validate the theoretical claim that entanglement and interference furnish a qualitatively new inductive bias for language understanding.

# Chapter 5

## Discussion

### 5.1 Synthesis of Theoretical and Empirical Evidence

The formal analysis in Chapter 3 predicts that projecting token embeddings into an entangled Hilbert space should linearize decision boundaries that would otherwise require wider hidden layers. The empirical study in Chapter 4 confirms this expectation: Quantum-Enhanced Transformers (QET) achieve gains of +3.4 percentage points on AG News and +0.3 % on IMDB while operating with roughly five percent fewer parameters than the strongest classical multi-head baseline. Extending the architecture from a single quantum head ( $\varphi = 0$ ) to four heads initialized with various phases  $\varphi_h = \pi/4 + h\pi/8$  further enlarges the span of the composite kernel and lifts IMDB accuracy by an additional 0.28 % at a modest nine-percent parameter increase—far smaller than the twenty-five-percent overhead required to obtain an equal uplift with classical multi-head attention. These findings substantiate the claim that entanglement and phase diversity provide a principled substitute for sheer network width.

### 5.2 Interpretability and Representational Geometry

Spectral analyses of the attention matrices reveal that QET exhibits a steeper eigenvalue decay than its classical counterpart, indicating that a small set of global modes captures the bulk of the relevance mass; this spectral coherence explains the strong topic-level recall observed on AG News. Complementary t-SNE projections show a twelve-percent reduction in intra-class dispersion on SST-5, with the Very Positive cluster separating most cleanly, mirroring the ten-point F1 improvement reported for that class. These geometric signatures align with the kernel-linearization property and collectively demonstrate that quantum attention reorganizes the latent space in a way that benefits long-range semantic discrimination.

### 5.3 Entangling Ansatz Architecture

The model employs a fixed two-layer, six-qubit, strongly entangling ansatz composed of single-qubit Pauli rotations followed by pairwise CNOT chains. This design delivers three concrete advantages.

**First, expressivity per parameter:** the exponential Hilbert dimension of a six-qubit register ( $2^6 = 64$ ) allows the kernel  $\Phi_\theta$  to represent higher-order token interactions that

would otherwise require wider classical layers, thereby underpinning the five-percent parameter savings reported in Chapter 4.

**Second, trainability:** two circuit layers keep the parameter surface shallow enough to avoid barren-plateau vanishing gradients, while the CNOT pattern guarantees full-qubit connectivity so that every rotation can influence every output amplitude within two timesteps.

**Third, hardware realism:** the depth–width trade-off respects the gate depth limits of near-term devices; a depth-two circuit executes well below typical coherence windows, and the uniform CNOT topology facilitates mapping onto common superconducting and trapped-ion layouts. Together these properties ensure that the ansatz offers a balanced compromise between expressivity, optimization stability and practical deployability, forming the backbone of the empirical gains demonstrated throughout this study.

## 5.4 Comparative Perspective

Relative to fixed random-feature quantum kernels, the end-to-end-trained map  $\Phi_\theta$  used here secures an additional 1.7 % on SST-5 under identical depth and qubit budgets, and its multi-phase design avoids the barren-plateau effect by maintaining informative gradients across all layers. Compared with efficient classical attention mechanisms that rely on low-rank or convolutional approximations and excel only at very long sequences ( $L > 2k$ ), QET is optimized for the mid-length regime ( $32 \leq L \leq 1024$ ) where semantic diversity is high and locality assumptions start to fail, thereby filling an important performance gap in current model families.

## 5.5 Limitations

The present implementation employs a fixed two-layer, six-qubit circuit that under-expresses deeply nested syntactic constructs and contributes to the 2.1 % deficit observed on SST-2. All experiments rely on tensor-network simulation; preliminary hardware trials with a depolarizing noise rate of 0.1 % reduce accuracy by 0.6 pp, indicating that some performance will be sacrificed on near-term devices. Memory usage remains quadratic in sequence length because the attention matrix is dense, so extremely long texts will require either sparse kernel approximations or blockwise processing to keep the model tractable.

## 5.6 Broader Impact and Concluding Remarks

Hybrid quantum NLP promises to reduce the parameter-to-performance ratio of state-of-the-art language models, but lower parameter counts do not automatically translate into wider hardware accessibility or lower environmental cost. Responsible deployment will require open-sourcing noise-aware implementations, publishing transparent carbon metrics, and rigorously auditing potential biases introduced by phase-controlled interference patterns. With these caveats in mind, the present work demonstrates that entanglement-aware kernels and phase diversity offer a viable route to parameter-efficient language understanding, establish a concrete empirical baseline, and provide an analytical toolkit for exploring the quantum–classical frontier in NLP.

# Chapter 6

## Future Directions

This chapter outlines concrete research avenues that build on the results presented thus far. Each subsection discusses a single axis of innovation—ansatz optimization, network complexity, task-specific design, and data-centric tuning, articulating both the scientific motivation and the practical steps required to realize further gains.

### 6.1 Optimising the Circuit Ansatz

The two-layer, six-qubit ansatz has proved sufficient for medium-length language tasks, yet it leaves accuracy headroom on syntax-dense corpora. A natural extension is *depth-adaptive entanglement*: increase the number of rotation–CNOT layers for sentences whose parse trees exceed a given depth threshold, thus tailoring expressive power to syntactic complexity without penalizing simpler inputs. In parallel, parameter sharing across layers can curb the depth-induced growth in gate count, maintaining hardware feasibility while enlarging the effective Hilbert space explored during training.

### 6.2 Expanding Network Complexity

While the current architecture mirrors the classical Transformer’s token-wise feed-forward block, more elaborate quantum–classical hybrids could interleave *quantum bottleneck layers* with conventional self-attention, allowing the model to alternate between global, entanglement-rich transformations and local, high-resolution updates. Another promising path involves *hierarchical kernels*: stack low- and high-resolution quantum heads so that shallow circuits capture surface-level correlations and deeper circuits refine class boundaries in a coarse-to-fine schedule, emulating the success of multiscale vision transformers in computer vision.

### 6.3 Task-Specific Ansatz and Network Design

Certain applications, legal reasoning, protein sequence analysis, code completion, exhibit domain constraints that can be embedded directly into the circuit topology. For example, a reversible ansatz that preserves token order may benefit algorithmic reasoning, while a symmetry-constrained entanglement pattern could exploit the permutation invariance of molecular graphs. Aligning the inductive bias of the circuit with the domain structure

is expected to offset the qubit overhead that generic architectures incur when forced to learn such constraints implicitly.

## 6.4 Data-Centric Quantum Tuning

Beyond architectural changes, the quality of the quantum feature map can be improved through *data-reweighted training*. Reweighting schemes that accentuate samples with high kernel-induced margin but low soft-max confidence can encourage the model to allocate quantum capacity where classical heads are weak. Curriculum schedules that present increasingly long or syntactically complex sentences only after the circuit parameters stabilize in simpler examples may further reduce gradient noise, shortening convergence time, and mitigating barren-plateau risk even in deeper circuits.

Taken together, these directions suggest a roadmap for closing the performance gap on phrase-level sentiment, scaling to documents longer than two thousand tokens and deploying on noisy hardware without sacrificing accuracy. By co-designing ansatz depth, network topology, and data curricula, future quantum-classical hybrids can transcend the limitations of current models, advancing both practical NLP performance and our theoretical understanding of quantum inductive biases.

# Chapter 7

## Conclusion

This study presents the Quantum-Enhanced Transformer (QET), a hybrid architecture that injects trainable quantum kernels and phase-controlled interference into the attention mechanism of classical Transformers. Across four representative NLP benchmarks, AG News, IMDB, SST-2 and SST-5, QET attains state-of-the-art or comparable accuracy while operating with approximately five percent fewer parameters than its strongest classical counterpart. On topic classification QET secures a decisive +3.4 % gain, achieves parity on long-form sentiment, and remains within 2 % on phrase-level sentiment despite its leaner footprint. Extending the design to four quantum heads widens the composite kernel’s span and delivers an additional 0.28 % on IMDB for a modest +9.7 % parameter overhead, less than half the cost required by a classical multi-head Transformer to obtain the same uplift.

Qualitative diagnostics corroborate the theoretical claims. Attention spectra reveal flatter, globally coherent similarity maps, and t-SNE projections exhibit tighter intra-class clusters and wider inter-class gaps, indicating that quantum entanglement reorganizes the latent space into a geometry that better separates semantic categories. The two-layer, six-qubit strongly entangling ansatz balances expressivity, trainability and hardware realism: it captures higher-order token interactions that would otherwise necessitate a wider network, avoids barren-plateau gradients, and respects the gate-depth limits of current NISQ devices.

Despite these advances, three limitations remain. First, the fixed ansatz depth under-expresses deeply nested syntax, leaving a two-percentage-point deficit on SST-2. Second, all experiments are simulator-based; preliminary shots on noisy hardware indicate a small but measurable accuracy drop, underscoring the need for error mitigation. Third, dense attention retains  $\mathcal{O}(L^2)$  memory cost, so extremely long documents will require sparse or blockwise quantum kernels.

Looking ahead, depth-adaptive entanglement, polarity-aware phase gating, and cross-modal quantum kernels offer promising routes to close the gap on short-sentence sentiment and to scale beyond two-thousand-token contexts. By demonstrating that entanglement-aware kernels and phase diversity yield parameter-efficient gains *today*, this work establishes a solid empirical baseline and an analytical toolkit to advance the quantum–classical frontier in natural-language understanding.

# Bibliography

- [1] Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., and Polosukhin, I. “Attention is All You Need.” *Advances in Neural Information Processing Systems*, 2017.
- [2] Devlin, J., Chang, M. W., Lee, K., and Toutanova, K. “BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding.” *NAACL-HLT*, 2019.
- [3] Schuld, M., Sinayskiy, I., and Petruccione, F. “An Introduction to Quantum Machine Learning.” *Contemporary Physics*, 2015.
- [4] Cerezo, M., Arrasmith, A., Babbush, R., Benjamin, S. C., Endo, S., Fujii, K., McClean, J. R., Mitarai, K., Yuan, X., Cincio, L., and Coles, P. J. “Variational Quantum Algorithms.” *Nature Reviews Physics*, 2021.
- [5] Biamonte, J., Wittek, P., Pancotti, N., Rebentrost, P., Wiebe, N., and Lloyd, S. “Quantum Machine Learning.” *Nature*, 2017.
- [6] Schuld, M., Bocharov, A., Svore, K., and Wiebe, N. “Supervised Learning with Quantum Computers.” *Nature*, 2019.
- [7] Khatri, N., Matos, G., Coopmans, L., and Clark, S. “Quixer: A Quantum Transformer Model.” *arXiv preprint, arXiv:2104.00756*, 2021.
- [8] Widdows, D., Karan, M., and Patel, R. “Quantum Natural Language Processing: A Survey.” *Springer*, 2024.
- [9] Mari, A., Bromley, T. R., Izaac, J., Schuld, M., and Killoran, N. “Transfer Learning in Hybrid Classical-Quantum Neural Networks.” *Physical Review Letters*, 2020.
- [10] Dey, P., Kumar, A., and Singh, T. “Hybrid Classical-Quantum Transfer Learning in Text Classification.” *Springer*, 2024.
- [11] Cong, I., Choi, S., and Lukin, M. D. “Quantum Convolutional Neural Networks.” *Nature Physics*, 2019.
- [12] Gao, L., Zhang, J., and Wang, Y. “Fast Quantum Algorithms for Attention Computation.” *arXiv preprint, arXiv:2307.08045*, 2023.
- [13] Chen, S. Y. C., Yoo, S., and Fang, Y. L. “Quantum Long Short-Term Memory.” *ICASSP*, 2022.
- [14] Li, G., Zhao, X., and Wang, X. “Quantum Self-Attention Neural Networks.” *Science China Information Sciences*, 2024.

- [15] O’Riordan, J., Green, T., and Kaur, P. “Hybrid Quantum-Classical Workflows for NLP.” arXiv preprint, arXiv:2004.06800, 2020.
- [16] Dong, D., Chen, C., Li, H., and Tarn, T. J. “Quantum Reinforcement Learning.” IEEE Transactions on Systems, Man, and Cybernetics, 2008.
- [17] M. Schuld, A. Bocharov, K. Svore, and N. Wiebe, “Circuit-centric quantum classifiers,” *Physical Review A*, vol. 101, no. 3, p. 032308, 2020.
- [18] Cherrat, El Amine, Iordanis Kerenidis, Natansh Mathur, Jonas Landman, Martin Strahm, and Yun Yvonna Li. *Quantum vision transformers*. arXiv preprint arXiv:2209.08167, 2022.