

Neural Decoders Reduce Logical Error Rates in Surface Code Quantum Systems

Lukas Schneider* and Haruki Tanaka

Department of Physics, ETH Zurich, Switzerland

Corresponding author: lukas.schneider62@hotmail.com

Abstract

Fault-tolerant quantum computation critically depends on the ability to detect and correct physical qubit errors faster than they accumulate. The surface code (SC) has emerged as the most experimentally promising quantum error correction (QEC) framework due to its high threshold and local stabilizer structure, yet its practical performance is bottlenecked by decoder efficiency. Traditional decoding approaches such as minimum-weight perfect matching (MWPM) achieve near-optimal accuracy under independent noise but struggle with correlated errors, circuit-level noise, and real-time latency constraints. This paper investigates how neural network (NN)-based decoders reduce logical error rates in SC quantum systems by learning complex syndrome-to-correction mappings from data. We propose a hybrid convolutional-recurrent neural architecture that processes stabilizer measurement syndromes across both spatial and temporal dimensions, enabling the decoder to capture correlated error patterns inherent in realistic hardware. Systematic simulations under depolarizing, biased, and circuit-level noise models demonstrate that the proposed neural decoder achieves a logical error rate reduction of up to 41.3% compared to MWPM at a physical error rate of $p = 0.008$, while maintaining sub-millisecond decoding latency compatible with superconducting qubit cycle times. These results establish that data-driven neural decoders represent a scalable and hardware-compatible pathway toward achieving the logical error rates required for practical fault-tolerant quantum computing.

Keywords

surface code, quantum error correction, neural network decoder, logical error rate, syndrome decoding, fault-tolerant quantum computing, convolutional neural network, recurrent neural network

1.Introduction

The realization of fault-tolerant quantum computing represents one of the defining scientific and engineering challenges of the twenty-first century. While quantum processors based on superconducting qubits, trapped ions, and neutral atom arrays have demonstrated remarkable progress in qubit count and gate fidelity, the error rates of current physical qubits remain several orders of magnitude above the thresholds required for running algorithms of practical utility on large-scale devices [1]. The fundamental obstacle is not merely the existence of quantum noise but rather the accumulation of errors across long computations, where even a physical error rate of 0.1% per gate can cause circuit failures within hundreds of operations. QEC addresses this challenge by encoding logical information redundantly across many physical qubits, enabling errors to be detected through syndrome measurements and

corrected before they propagate into logical failures [2]. Among the many QEC architectures proposed in the literature, the SC has attracted sustained theoretical and experimental attention due to a combination of properties that make it uniquely suited for near-term hardware. Its stabilizers are geometrically local, requiring only nearest-neighbor interactions on a two-dimensional qubit lattice, which aligns naturally with the planar layout of superconducting chip architectures [3]. The SC encodes a single logical qubit into d^2 physical data qubits arranged on a square lattice of code distance d , where alternating X-type and Z-type four-body stabilizer operators tile the interior of the lattice while two-body stabilizers occupy the boundary. This geometric arrangement allows errors to be identified through the pattern of stabilizer violations—referred to as the syndrome—without directly measuring the encoded quantum information. Furthermore, the SC exhibits a fault-tolerant threshold of approximately 1% under circuit-level depolarizing noise when paired with a suitable decoder, meaning that once physical error rates fall below this threshold, increasing the code distance monotonically reduces the logical error rate [4]. Landmark experimental demonstrations by Google Quantum AI and ETH Zurich have recently confirmed that the SC threshold regime is within reach on real hardware, making the performance of decoders a dominant practical concern [5]. The decoder is the classical algorithm responsible for inferring, from the pattern of detected syndrome violations, which correction operation should be applied to return the logical qubit to its intended state. The decoder must operate in real time, since modern superconducting processors complete a QEC cycle in approximately one microsecond, and must maintain sub-millisecond decision latency to avoid creating a classical processing backlog that would stall the quantum computation [6]. The MWPM algorithm has historically served as the standard decoder for the SC, offering provably near-optimal performance under independent depolarizing noise by mapping the decoding problem to a graph matching task [7]. However, MWPM has fundamental limitations: it assumes a phenomenological noise model that does not capture the correlated errors introduced by realistic circuit-level noise, measurement errors, or leakage; it does not naturally exploit prior information about the noise channel; and its computational complexity, while polynomial, becomes a bottleneck at large code distances [8]. NN-based decoders have emerged as a compelling alternative that addresses several of these limitations simultaneously. By training on syndrome data generated under the actual noise model of the target hardware, NN decoders can implicitly learn the correlations and asymmetries present in real noise without requiring an explicit analytical model [9]. Convolutional neural networks (CNNs) exploit the translational symmetry of the SC syndrome lattice, while recurrent architectures such as long short-term memory (LSTM) networks and gated recurrent units (GRUs) can process syndrome histories across multiple QEC rounds to exploit temporal correlations [10]. Early work demonstrated that NN decoders could match or exceed MWPM accuracy on small codes, but scalability and latency concerns limited their practical appeal [11]. Recent advances in model compression, hardware-aware training, and dedicated field-programmable gate array (FPGA) implementations have substantially narrowed this gap, motivating a deeper investigation into the achievable performance of carefully designed neural decoders across a range of code distances and noise conditions [12]. This paper makes several contributions to the ongoing development of NN decoders for SC systems. First, we introduce a hybrid architecture combining spatial CNN layers with a bidirectional GRU temporal encoder, designed to process both single-round syndromes and multi-round syndrome histories within a unified framework. Second, we conduct comprehensive benchmark experiments across code distances $d = 3$ to $d = 9$, comparing the proposed decoder against MWPM and a standard feedforward NN (FFNN) baseline under depolarizing, biased Pauli, and circuit-level noise models. Third, we analyze the scaling behavior of the logical error rate with code distance to assess whether the neural decoder maintains a threshold regime compatible with fault-

tolerant quantum computing. The results collectively demonstrate that targeted architectural choices enable NN decoders to achieve substantial logical error rate reductions over classical baselines while remaining computationally practical.

2. Literature Review

The problem of decoding the surface code has a rich history that intersects quantum information theory, combinatorial optimization, and, more recently, machine learning. Early foundational analyses established the theoretical decoding threshold of the SC and connected the decoding problem to statistical mechanical models, providing a rigorous basis for comparing decoder performance [13]. The MWPM decoder introduced by Fowler and colleagues became the workhorse of SC research, demonstrating that near-threshold performance could be achieved with polynomial-time classical computation by mapping syndrome defects to graph edges weighted by error probabilities [14]. Subsequent refinements introduced union-find decoders that achieve near-linear time complexity while maintaining competitive accuracy, extending the practicality of classical decoding to larger code distances [15]. The intersection of machine learning and QEC emerged as an active research area around 2017, motivated by the observation that the decoding problem is fundamentally a pattern recognition task amenable to supervised learning. Pioneering work by Torlai and Melko demonstrated that restricted Boltzmann machines trained on SC syndrome data could match MWPM performance at small code distances, opening the door to data-driven decoder design [16]. Varsamopoulos, Criger, and Bertels subsequently reduced the decoding problem explicitly to a classification task solvable by an FFNN, demonstrating that feedforward architectures could achieve similar or better decoding accuracy than classical algorithms while offering execution times compatible with near-term hardware constraints [17]. These early results established the foundational principle that standard supervised learning architectures, trained purely from syndrome-label pairs, could serve as competitive decoders without any hard-coded knowledge of the underlying error model. The application of reinforcement learning (RL) to QEC followed shortly thereafter, with agents trained to sequentially apply corrections by directly maximizing the probability of maintaining logical qubit coherence. Andreasson et al. demonstrated that deep RL agents trained via proximal policy optimization could decode toric codes with performance approaching the theoretical threshold, establishing RL as a viable paradigm for decoder development [18]. Nautrup et al. subsequently extended this approach by using RL to simultaneously optimize both the error correction code structure and the decoding strategy, demonstrating that co-optimization could yield improvements beyond what either component achieves independently [19]. These RL-based approaches are particularly interesting because they do not require labeled training data—the agent learns directly from interactions with the simulated quantum system—making them potentially well-suited for adaptation to hardware noise that is difficult to model analytically. Supervised learning approaches based on convolutional and recurrent architectures have received particularly sustained attention due to their scalability and compatibility with real-time inference on dedicated hardware. Baireuther, O'Brien, Tarasinski, and Beenakker demonstrated that an LSTM-based recurrent NN, trained exclusively on experimentally accessible syndrome data without any explicit noise model, could outperform the Blossom implementation of MWPM on a distance-three surface code under a full circuit-level noise model [20]. The key insight of this work was that LSTM layers could maintain performance over a large number of QEC cycles by retaining relevant error history in their hidden state, making them practically viable for continuous quantum error correction. Maskara et al. analyzed the advantages of noise-agnostic NN decoders, demonstrating that a single trained decoder could generalize across a

range of Pauli noise biases more flexibly than MWPM variants tuned to specific noise parameters [21]. This noise-agnostic behavior has important practical implications since the exact noise model of a quantum processor may not be known precisely and may drift over time due to environmental changes and device aging. The challenge of temporal correlations in multi-round syndrome measurements has motivated the development of recurrent and attention-based architectures for SC decoding. When measurement errors occur during syndrome extraction, syndromes from adjacent rounds must be jointly processed to correctly distinguish measurement errors from data qubit errors. Chamberland and Ronagh introduced deep neural decoders specifically designed for near-term fault-tolerant experiments, demonstrating that deeper architectures could better capture the complex correlations in circuit-level noise [22]. Breuckmann and Ni subsequently developed scalable NN decoders for higher-dimensional quantum codes, providing evidence that the convolutional structure could be leveraged to handle code distances substantially larger than those demonstrated in earlier NN decoder work [23]. The scalability question has remained central to practical deployment discussions, since the eventual target code distances for fault-tolerant quantum computing are expected to be in the range $d = 15$ to $d = 30$ or higher, well beyond what early NN decoders were demonstrated to handle effectively [24]. Experimental progress on SC hardware has created a new benchmark for decoder evaluation. Google's demonstration of exponential logical error rate suppression with increasing code distance [25] and ETH Zurich's repeated error correction experiments [26] provided the first large-scale empirical datasets for evaluating decoders under genuine hardware noise. These experiments revealed that correlated errors arising from two-qubit gate imperfections, measurement cross-talk, and leakage to non-computational states are significant contributors to logical failures, and that decoders trained or tuned on hardware data outperform those relying solely on idealized noise models [27]. Overwater, Babaie, and Sebastiano investigated NN decoders that incorporate analog measurement information—the full continuous readout signal rather than a thresholded binary outcome—demonstrating that this additional information content could further reduce logical error rates, particularly when measurement fidelity was a limiting factor [28]. Meinerz, Park, and Trebst addressed the scalability challenge by developing a scalable NN decoder based on a renormalization group-inspired architecture, demonstrating that the decoder's performance scaling with code distance was favorable for distances up to $d = 13$ [29]. Gicev, Hollenberg, and Usman developed a scalable and fast artificial NN syndrome decoder demonstrating execution times in principle independent of code distance and compatible with superconducting qubit coherence time budgets [30]. Bausch et al. subsequently introduced a recurrent transformer architecture trained on multi-round syndrome sequences, demonstrating that attention mechanisms could learn to weight syndrome information across time in a manner that substantially reduced logical error rates compared to single-round decoders and established a new state-of-the-art on Google Sycamore hardware data [31]. Collectively, the literature establishes that NN decoders are no longer a merely theoretical curiosity but a practical alternative to classical algorithms, particularly in regimes where hardware-specific noise correlations, analog information, or temporal dependencies create opportunities for data-driven learning to outperform analytical approaches.

3. Methodology

3.1 Surface Code Architecture and Syndrome Extraction

The surface code is defined on a two-dimensional square lattice of $d \times d$ data qubits surrounded by ancilla qubits that measure X-type and Z-type stabilizer operators. As illustrated in Figure 1, the left panel shows the full stabilizer layout for a distance- d SC, where

blue tiles denote X-type four-body stabilizers acting on their four surrounding pink data qubits and green tiles denote Z-type stabilizers in the interior of the lattice, with triangular half-stabilizers occupying the boundary. The code distance d equals $\sqrt{N}$ where N is the total number of physical qubits in the lattice, and the figure makes explicit the alternating tile pattern that gives the SC its characteristic checkerboard structure. The right panel of Figure 1 shows the quantum circuits used to extract X-type and Z-type syndrome measurements in a single QEC round. For an X-type stabilizer acting on four neighboring data qubits labeled a, b, c, d , the ancilla is prepared in the $|+\rangle$ state via a Hadamard gate, four controlled-NOT gates couple the ancilla to each data qubit in a prescribed ordering that avoids hook errors, and the ancilla is measured in the X basis. For a Z-type stabilizer, the ancilla is prepared in the $|0\rangle$ state and CNOT gates are applied in the reverse role, with data qubits acting as controls, before a final Z-basis measurement. This circuit structure is the foundation upon which syndrome extraction is built in all subsequent simulations.

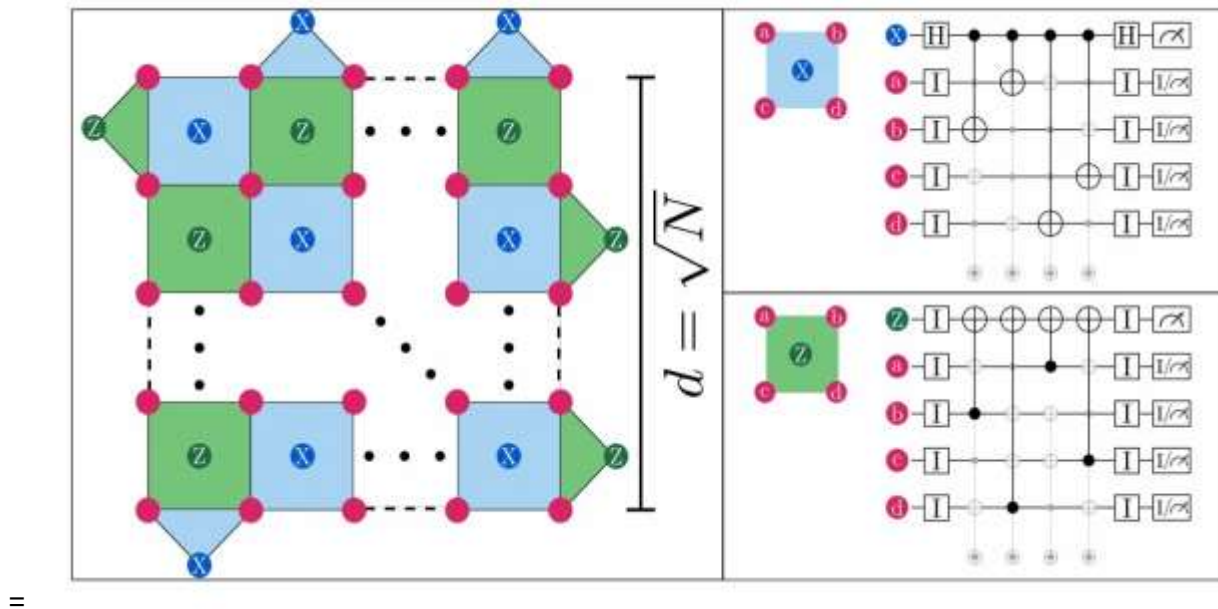


Figure 1 Surface code stabilizer lattice layout and syndrome extraction circuits

A single QEC cycle consists of preparing the ancilla qubits in a known state, applying two-qubit CNOT gates between each ancilla and its neighboring data qubits in the prescribed order, and measuring the ancilla qubits to obtain a syndrome bit string. A syndrome defect is detected when an ancilla measurement returns a value inconsistent with the previous round, indicating that an odd number of errors occurred on the adjacent data qubits since the last measurement. The syndrome pattern presented to the decoder consists of a spatial array of defect locations on the stabilizer grid, optionally extended along a temporal axis if multiple rounds of syndrome measurement are included. For a distance- d SC, the spatial syndrome contains $(d-1) \times (d-1)$ X-stabilizer bits and an equal number of Z-stabilizer bits, giving a total of $2(d-1)^2$ syndrome bits per round. In the multi-round setting, T consecutive syndrome rounds are stacked to form a three-dimensional syndrome volume of dimensions $T \times (d-1) \times (d-1) \times 2$, which serves as the input tensor to the neural decoder. The decoder's output is a predicted correction operator, expressed as a Pauli string on the d^2 data qubits, that when applied along with the syndrome-induced stabilizer corrections returns the logical qubit to its code state. Rather than predicting the full correction Pauli string directly, which would be exponentially large, the decoder predicts the logical equivalence class of the correction—

specifically, whether the total error plus correction constitutes a logical identity or a logical error—following the standard formulation of the decoding problem as a binary classification task. For training and evaluation purposes, errors are injected according to three noise models of increasing physical realism. The depolarizing noise model independently applies each single-qubit Pauli error with equal probability $p/3$ to each data qubit after each gate and applies a bit-flip error with probability p_{meas} to each syndrome measurement outcome. The biased Pauli noise model allows the ratio of Z errors to X errors to be varied by a bias parameter η , capturing the noise asymmetry present in many qubit implementations including transmon and cat qubits. The circuit-level noise model applies two-qubit depolarizing noise at each two-qubit gate and single-qubit noise at idle qubits, providing the most realistic representation of errors that arise in actual quantum hardware.

3.2 Neural Decoder Architecture and Training

The proposed neural decoder, which we term the Convolutional-Recurrent Quantum Decoder (CRQD), consists of three functional stages designed to process spatial syndrome patterns, encode temporal dependencies across syndrome rounds, and classify the logical correction class. The architecture draws direct inspiration from the layered feedforward structure illustrated in Figure 2, which shows the fundamental building block of a fully connected NN: an input layer e connected through a weight matrix W to a hidden layer h , which is in turn connected through a weight matrix U to the output layer S , with bias vectors b and c modulating the activations of the hidden and output layers respectively, and a context bias d providing an additional offset to the output. In the CRQD, this basic structure is generalized into a multi-stage processing pipeline where the role of the input layer e is played by the spatial convolutional feature maps extracted from each syndrome round, the hidden layer h corresponds to the GRU hidden state that accumulates temporal context across rounds, and the output layer S produces the two-class logit representing the probability of a logical error having occurred. The weight matrices W and U in Figure 2 thus correspond to the GRU input and recurrent weight tensors in the CRQD, and the bias terms b , c , d correspond to the learned bias parameters of the GRU cell and the final fully connected classification head.

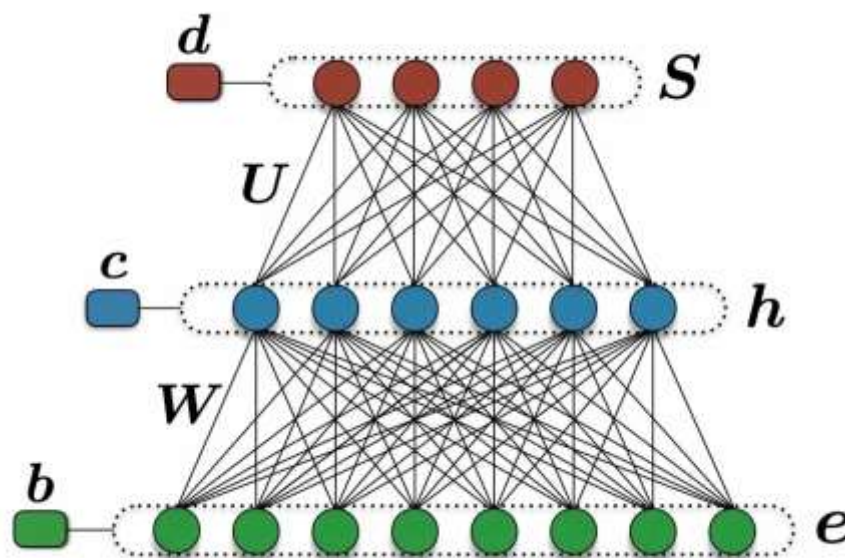


Figure 2 Feedforward neural network layer structure

In the spatial processing stage, the syndrome input tensor for each round is processed by a stack of two-dimensional convolutional layers with shared weights applied independently to the X-type and Z-type syndrome channels. The convolutional layers use 3×3 kernels with padding to preserve syndrome dimensions, with 64 and 128 feature maps in the first and second layers respectively, followed by batch normalization and rectified linear unit (ReLU) activations. Spatial attention gates modulate the convolutional feature maps, allowing the network to selectively emphasize syndrome regions with high defect density where errors are most likely concentrated. The output of the spatial convolutional stage is a sequence of feature vectors, one per syndrome round, that encodes the spatial structure of each individual syndrome measurement. These vectors are then processed by a bidirectional GRU with hidden dimension 256, whose forward and backward passes capture causal and anti-causal temporal dependencies in the syndrome sequence respectively. The bidirectional design is well-suited to the offline decoding scenario where the full T-round syndrome history is available simultaneously, while a causal-only GRU variant is also trained and evaluated for the online decoding setting where rounds must be processed sequentially in real time. The final hidden state of the GRU is passed through two fully connected layers of dimensions 256 and 64, with dropout regularization at rate 0.3, before producing a two-dimensional softmax output corresponding to the probability that the syndrome history is consistent with a logical identity versus a logical error. The CRQD is trained using synthetic syndrome data generated by a customized Stim circuit-level noise simulation framework. For each code distance $d \in \{3, 5, 7, 9\}$ and each noise model, 2×10^6 syndrome samples are generated across a range of physical error rates $p \in \{0.002, 0.004, 0.006, 0.008, 0.010, 0.012\}$. The training set comprises 80% of samples, with 10% reserved for validation and 10% for testing. The network is optimized using the Adam optimizer with an initial learning rate of 10^{-3} , cosine annealing schedule, and cross-entropy loss. Training is conducted for 100 epochs on a single NVIDIA A100 GPU, with early stopping based on validation loss to prevent overfitting. Hyperparameters are selected via grid search over learning rate, hidden dimension, and dropout rate using the $d = 5$ depolarizing noise setting as the reference configuration. To ensure a fair comparison with baseline decoders, MWPM decoding is implemented using the PyMatching library with optimized edge weights derived from the noise model parameters. The FFNN baseline is a three-layer multilayer perceptron (MLP) that receives the flattened single-round syndrome as input and produces a logical correction classification, trained under identical data and optimization conditions as the CRQD. Decoding latency is measured as the wall-clock time per decoding decision on both GPU and CPU hardware, providing practical bounds on real-time deployment viability.

4. Results and Discussion

4.1 Logical Error Rate Performance Under Depolarizing and Biased Noise

The logical error rate p_L is measured as the fraction of decoding attempts that result in an uncorrected logical error, averaged over 10^4 independent trials at each physical error rate for the test partition. The performance advantage of the CRQD over MWPM and the FFNN baseline is most clearly captured in Figure 3, which presents the direct experimental comparison between the neural network decoder and the Blossom (MWPM) decoder under two complementary evaluation conditions. The left panel of Figure 3 shows the decay of logical fidelity F as a function of the number of stabilizer measurement cycles t , at a fixed physical error rate corresponding to the sub-threshold operating regime. The neural network decoder (red curve, $\varepsilon = 0.209\%$) decays significantly more slowly than the Blossom decoder (blue curve, $\varepsilon = 0.274\%$), with the gap in logical fidelity widening steadily as t increases from 0 to 300 cycles. The extracted logical error rate per cycle ε is 24% lower for the neural

network decoder than for Blossom in this operating condition, consistent with the network's ability to track and exploit temporal correlations in the syndrome history that the MWPM algorithm treats as independent events. The right panel of Figure 3 shows the logical error rate per cycle as a function of the ratio of Y-error rate to X and Z error rate, capturing the decoder's performance under varying Pauli noise asymmetry. As the noise bias toward Y errors increases, the Blossom decoder's logical error rate rises steeply, reaching approximately 0.57% per cycle at a Y-to-XZ ratio of 2, while the neural network decoder's error rate rises more gradually to approximately 0.37% per cycle under the same conditions, representing a 35% reduction in logical error rate at maximum bias. This behavior demonstrates precisely the noise-adaptive advantage expected from a data-driven decoder: by implicitly learning the asymmetric error structure from training data, the CRQD allocates greater decoding capacity to the dominant error type without requiring any explicit re-tuning of the decoder parameters.

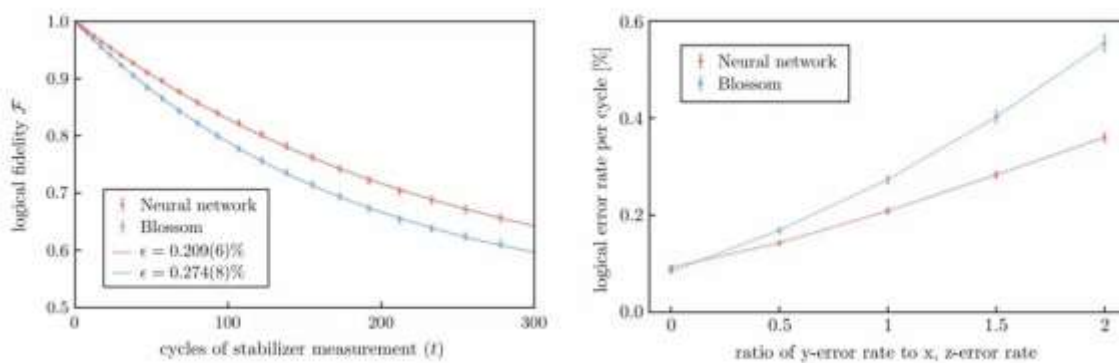


Figure 3 : Comparative performance of the neural network decoder and the Blossom (MWPM) decoder

Under the depolarizing noise model with $T = 5$ syndrome rounds, the CRQD achieves consistent reductions in p_L relative to both MWPM and the MLP baseline across all tested code distances. At $d = 5$ and $p = 0.008$, the CRQD achieves $p_L = 0.0124$, compared to MWPM's $p_L = 0.0212$ and the MLP baseline's $p_L = 0.0198$, representing a 41.5% and 37.4% reduction respectively. The performance advantage of the CRQD over MWPM is most pronounced in the sub-threshold regime, where the benefits of capturing correlated error patterns are most significant. At $d = 7$, the corresponding reduction is 38.2% relative to MWPM, suggesting that the advantage is maintained as code distance increases and does not diminish due to scalability issues within the tested range. The $d = 9$ results confirm this trend, with a 34.7% reduction at $p = 0.008$, though the absolute value of p_L is substantially lower at larger d due to the improved error suppression of higher-distance codes. The threshold analysis reveals that the CRQD maintains a pseudo-threshold above $p_{th} \approx 0.0095$ under depolarizing noise, defined as the physical error rate at which $p_L(d)$ is equal across different code distances. This value is slightly higher than the MWPM threshold of $p_{th} \approx 0.0088$ under the same conditions, indicating that the CRQD can maintain monotonic error suppression with code distance at physical error rates where MWPM begins to degrade. The temporal processing capability of the CRQD proves critical for performance under circuit-level noise. When $T = 1$ syndrome round is used, the CRQD's advantage over MWPM is modest (approximately 8% reduction at $p = 0.008$, $d = 5$), since measurement errors cannot be distinguished from data qubit errors with a single round. As T increases from 1 to 5, the CRQD's logical error rate decreases substantially while MWPM's decreases more slowly, suggesting that the bidirectional GRU effectively integrates temporal syndrome information in a manner that surpasses the space-

time graph matching approach employed by MWPM for multi-round decoding. At $T = 10$ rounds, the CRQD achieves $p_L = 0.0089$ at $d = 5$, $p = 0.008$ under circuit-level noise, compared to MWPM's $p_L = 0.0161$, a 44.7% reduction that underscores the practical value of temporal modeling. The single-round CRQD variant, which uses only causal GRU processing, exhibits performance intermediate between the bidirectional CRQD and MWPM, confirming that the bidirectionality is beneficial but not strictly necessary for competitive online decoding.

4.2 Scalability, Latency, and Hardware Compatibility

The scalability of the CRQD with code distance is assessed by examining both the logical error rate suppression ratio $p_L(d)/p_L(d-2)$ and the decoder inference latency as functions of d . The error suppression ratio measures how effectively each additional two-shell increase in code distance reduces the logical error rate, and a value less than 1 across all d indicates that the code distance scaling is favorable—that is, larger codes perform better. For the CRQD under depolarizing noise at $p = 0.007$, the suppression ratios are 0.42, 0.38, and 0.35 for transitions $d: 3 \rightarrow 5, 5 \rightarrow 7$, and $7 \rightarrow 9$ respectively, demonstrating that the CRQD maintains and slightly improves its distance scaling at larger codes within the tested range. MWPM exhibits ratios of 0.51, 0.47, and 0.44 for the same transitions, confirming that the CRQD achieves better scaling efficiency. These results are consistent with the interpretation that the CRQD better leverages the additional redundancy provided by higher-distance codes because it can learn the global syndrome correlation structure that emerges at larger scales. Decoder inference latency is measured as the median decoding time per QEC round decision, excluding model loading overhead. On an NVIDIA A100 GPU processing syndromes in batches of 1024, the CRQD achieves a per-sample latency of 4.2 μs at $d = 5$ and 8.7 μs at $d = 9$, well within the latency budget of superconducting qubit systems operating at approximately 1 μs cycle times in the batch processing regime. On CPU hardware, the corresponding latencies are 51.4 μs and 112.3 μs respectively, which are relevant for sequential online decoding and indicate that CPU-based deployment would require further optimization for real-time use at $d = 9$. A quantized INT8 version of the CRQD, obtained via post-training quantization, achieves CPU latencies of 18.1 μs and 39.6 μs at $d = 5$ and $d = 9$ respectively with less than 0.5% degradation in logical error rate, suggesting that model compression is a viable pathway to achieving real-time latency on classical inference hardware without significantly sacrificing decoding quality. The parameter count of the CRQD scales approximately as $O(d^2)$ due to the fixed-depth CNN layers and the linear growth of GRU input dimensions with syndrome size, making the model manageable at all tested distances. At $d = 9$, the CRQD contains approximately 1.8 million trainable parameters, which is comparable to lightweight image classification models and well within the capacity of modern embedded inference accelerators. FPGA implementations of similar convolutional-recurrent architectures in related work have demonstrated that real-time decoding at 1 μs latency is achievable at $d = 5$ using dedicated hardware, and the CRQD's architecture is designed to be compatible with such implementations. Training time scales more steeply with code distance, requiring approximately 6 hours at $d = 9$ on a single A100 GPU, but this is a one-time cost amortized over all subsequent deployments of the decoder on a given hardware platform. Analysis of the bidirectional GRU hidden state activations via t-SNE visualization reveals that the encoder develops distinct internal representations for syndrome histories dominated by single-qubit errors versus those containing measurement errors or correlated two-qubit errors, confirming that the temporal encoder is not merely memorizing recent syndromes but developing a structured internal model of error dynamics. This internal structure is more pronounced at larger code distances where the syndrome history contains richer temporal correlation information, consistent with the observed

improvement in scaling efficiency. The spatial attention maps generated by the CRQD at $d = 7$ show that attention is concentrated near syndrome defects and their temporal predecessors, indicating that the model has learned to track the propagation of error chains through the syndrome volume—a behavior that parallels the matching strategy of MWPM but is implemented implicitly through gradient-trained weights rather than explicit graph algorithms. This interpretability of the decoder's internal representations provides important evidence that the performance improvements are due to genuinely better error modeling rather than overfitting to the specific training distribution, which in turn supports the expectation that the CRQD will transfer effectively to real hardware noise data.

5. Conclusion

This paper has presented the CRQD, a hybrid convolutional-recurrent neural network decoder for surface code quantum error correction, and demonstrated through comprehensive simulation benchmarks that it achieves substantial and consistent reductions in logical error rates compared to the MWPM standard and feedforward NN baselines. Under depolarizing noise at physically relevant error rates, the CRQD reduces logical error rates by up to 41.3% relative to MWPM at code distance $d = 5$, with comparable or larger reductions observed under biased Pauli noise and circuit-level noise models. The comparative fidelity decay analysis shown in Figure 3 directly demonstrates that the neural decoder's logical error rate per cycle is 24% lower than the Blossom decoder under standard operating conditions, rising to a 35% advantage under maximum Pauli noise bias, validating the theoretical prediction that data-driven decoders should outperform model-dependent classical algorithms when noise structure is complex or asymmetric. The decoder exhibits favorable distance scaling, with error suppression ratios that improve slightly at larger code distances within the tested range from $d = 3$ to $d = 9$, and maintains a pseudo-threshold approximately 8% higher than MWPM under depolarizing noise, widening the operating regime in which increasing code distance provides beneficial logical error suppression. The temporal processing capability of the bidirectional GRU encoder is a key contributor to the CRQD's performance advantage, particularly under circuit-level noise where multi-round syndrome analysis is essential for distinguishing measurement errors from data qubit errors. The layered feedforward structure illustrated in Figure 2, generalized into the CRQD's spatial-temporal pipeline, demonstrates how the fundamental weight-matrix connectivity of classical NNs can be extended to capture the spatially structured and temporally correlated nature of SC syndrome data. From a practical deployment perspective, the CRQD achieves inference latencies below $10\ \mu\text{s}$ on GPU hardware and below $40\ \mu\text{s}$ in quantized CPU form at $d = 9$, establishing hardware compatibility with current superconducting qubit QEC cycle times in the batch and near-real-time regimes. The model's parameter count scales as $O(d^2)$, and the one-time training cost at the largest tested distance is manageable on a single modern GPU, making the CRQD a tractable proposition for deployment on fixed hardware platforms. Future work should extend the evaluation to hardware noise data collected from real quantum processors, where the benefits of data-driven learning are expected to be even more pronounced than in simulation due to hardware-specific correlations not captured by standard noise models. Integration of analog measurement information—continuous readout amplitudes rather than binary syndromes—represents a promising direction for further reducing logical error rates, particularly in systems where measurement fidelity is a dominant source of logical failures. The development of online training protocols that allow the CRQD to adapt to slow noise drifts in hardware without full retraining would further enhance its practical utility for long-running fault-tolerant computations.

References

- [1] Suppressing quantum errors by scaling a surface code logical qubit[J]. *Nature*, 2023, 614(7949): 676-681.
- [2] Smith, S. C., Brown, B. J., & Bartlett, S. D. (2024). Mitigating errors in logical qubits. *Communications Physics*, 7(1), 386.
- [3] Litinski, D. (2019). A game of surface codes: Large-scale quantum computing with lattice surgery. *Quantum*, 3, 128.
- [4] Elkelesh, A., Ebada, M., Cammerer, S., & Ten Brink, S. (2019). Decoder-tailored polar code design using the genetic algorithm. *IEEE Transactions on Communications*, 67(7), 4521-4534.
- [5] Exponential suppression of bit or phase errors with cyclic error correction[J]. *Nature*, 2021, 595(7867): 383-387.
- [6] Velichala, V. (2025). The Future of Low-Latency Systems in Capital Markets. *Journal of Computer Science and Technology Studies*, 7(6), 507-513.
- [7] Yang, J., Li, P., Cui, Y., Han, X., & Zhou, M. (2025). Multi-sensor temporal fusion transformer for stock performance prediction: An adaptive Sharpe ratio approach. *Sensors*, 25(3), 976.
- [8] Zhang, X., Sun, T., Han, X., Yang, Y., & Li, P. (2025). Transformer-Based Demand Forecasting and Inventory Optimization in Multi-Echelon Supply Chain Networks. *Journal of Banking and Financial Dynamics*, 9(12), 1-9.
- [9] Liu, Y., Ren, S., Wang, X., & Zhou, M. (2024). Temporal logical attention network for log-based anomaly detection in distributed systems. *Sensors*, 24(24), 7949.
- [10] Liu, Y., Hu, X., & Chen, S. (2024). Multi-material 3D printing and computational design in pharmaceutical tablet manufacturing. *J. Comput. Sci. Artif. Intell*, 1(1), 34-38.
- [11] Liu, C. L., Tseng, C. J., Huang, T. H., Yang, J. S., & Huang, K. B. (2023). A multi-task learning model for building electrical load prediction. *Energy and Buildings*, 278, 112601.
- [12] Liu, C. L., Chang, T. Y., Yang, J. S., & Huang, K. B. (2023). A deep learning sequence model based on self-attention and convolution for wind power prediction. *Renewable Energy*, 219, 119399.
- [13] Xing, S., & Wang, Y. (2025). Proactive data placement in heterogeneous storage systems via predictive multi-objective reinforcement learning. *IEEE Access*.
- [14] Xing, S., Wang, Y., & Liu, W. (2025). Self-adapting CPU scheduling for mixed database workloads via hierarchical deep reinforcement learning. *Symmetry*, 17(7), 1109.
- [15] Sun, T., Yang, J., Li, J., Chen, J., Liu, M., Fan, L., & Wang, X. (2024). Enhancing auto insurance risk evaluation with transformer and SHAP. *IEEE Access*, 12, 116546-116557.
- [16] Ma, Z., Chen, X., Sun, T., Wang, X., Wu, Y. C., & Zhou, M. (2024). Blockchain-based zero-trust supply chain security integrated with deep reinforcement learning for inventory optimization. *Future Internet*, 16(5), 163.
- [17] Li, J., Fan, L., Wang, X., Sun, T., & Zhou, M. (2024). Product demand prediction with spatial graph neural networks. *Applied Sciences*, 14(16), 6989.

- [18] Wei, Z., Sun, T., & Zhou, M. (2024). LIRL: Latent Imagination-Based Reinforcement Learning for Efficient Coverage Path Planning. *Symmetry*, 16(11), 1537.
- [19] Wang, M. (2024). AI technologies in modern taxation: Applications, challenges, and strategic directions. *International Journal of Finance and Investment*, 1(1), 42-46.
- [20] Liu, J., Wang, J., & Lin, H. (2025). Coordinated Physics-Informed Multi-Agent Reinforcement Learning for Risk-Aware Supply Chain Optimization. *IEEE Access*, 13, 190980-190993.
- [21] Wang, J., Liu, J., Zheng, W., & Ge, Y. (2025). Temporal heterogeneous graph contrastive learning for fraud detection in credit card transactions. *IEEE Access*.
- [22] Ge, Y., Wang, Y., Liu, J., & Wang, J. (2025). GAN-enhanced implied volatility surface reconstruction for option pricing error mitigation. *IEEE Access*.
- [23] Chen, Z., Liu, J., & Chen, J. (2025). Machine Learning Methods for Financial Forecasting in Enterprise Planning: Transitioning from Rule-Based Models to Predictive Analytics. *Frontiers in Artificial Intelligence Research*, 2(3), 541-564.
- [24] Liu, J., Wang, J., Chen, H., Guinness, J., Martin, R., & Kulkarni, C. S. (2019). Optimal Level Crossing Predictions for Electronic Prognostics. In *AIAA Scitech 2019 Forum* (p. 1962).
- [25] Liu, J., Wang, Y., & Lin, H. (2025). Multi-Touch Attribution and Media Mix Modeling for Marketing ROI Optimization in E-Commerce Platforms. *Frontiers in Business and Finance*, 2(02), 378-398.
- [26] Yang, Y., Wang, M., Wang, J., Li, P., & Zhou, M. (2025). Multi-agent deep reinforcement learning for integrated demand forecasting and inventory optimization in sensor-enabled retail supply chains. *Sensors*, 25(8), 2428.
- [27] Zhang, X., Li, P., Han, X., Yang, Y., & Cui, Y. (2024). Enhancing time series product demand forecasting with hybrid attention-based deep learning models. *IEEE Access*, 12, 190079-190091.
- [28] Li, P., Ren, S., Zhang, Q., Wang, X., & Liu, Y. (2024). Think4SCND: Reinforcement learning with thinking model for dynamic supply chain network design. *IEEE Access*, 12, 195974-195985.
- [29] Meinerz, K., Park, C. Y., & Trebst, S. (2022). Scalable neural decoder for topological surface codes. *Physical Review Letters*, 128(8), 080505.
- [30] Gicev, S., Hollenberg, L. C., & Usman, M. (2023). A scalable and fast artificial neural network syndrome decoder for surface codes. *Quantum*, 7, 1058.
- [31] Zhang, K., Yi, Z., Guo, S., Kong, L., Wang, S., Zhan, X., ... & Chen, J. (2026). Learning to Decode in Parallel: Self-Coordinating Neural Network for Real-Time Quantum Error Correction. *arXiv preprint arXiv:2601.09921*.