

SPEP: Adaptive Resource Orchestration in Quantum Networks through Strategic Entanglement Pre-Positioning

Vineet Kumar Dwivedi

*Department of Information Technology
National Institute of Technology Raipur
vkdwivedi589@gmail.com*

Vivek Shukla

*Department of Information Technology
National Institute of Technology Raipur
viveks2992@gmail.com*

Chandrashekar Jatoth

*Department of Information Technology
National Institute of Technology Raipur
jchandrashekar.it@nitrr.ac.in*

Rajesh Doriya

*Department of Information Technology
National Institute of Technology Raipur
rajeshdoriya.it@nitrr.ac.in*

Rajkumar Buyya

*School of Computing and Information Systems
The University of Melbourne, Australia
rbuyya@unimelb.edu.au*

Abstract—Quantum networks promise unprecedented capabilities for secure communication and distributed quantum computing through entanglement distribution. However, the fragile nature of quantum entanglement and its exponential decay present significant challenges for efficient resource allocation. This paper introduces SPEP (Strategic Predictive Entanglement Pre-positioning), a novel framework that leverages machine learning techniques to anticipate communication demands and proactively establish entangled links. We implement and evaluate multiple prediction models including LSTM neural networks and ARIMA time series analysis, comparing them against traditional reactive and greedy routing approaches. Through comprehensive simulations across four distinct network topologies with over multiple simulation runs, our SPEP framework demonstrates significant improvement in success rates, 15% reduction in average latency, and 25% better resource utilization compared to baseline methods. The results provide strong evidence that machine learning-driven prediction can significantly enhance quantum network performance, particularly in networks with temporal demand patterns.

Index Terms—quantum networks, entanglement distribution, SPEP framework, resource allocation, LSTM, network optimization

I. INTRODUCTION

The emergence of quantum networking represents a paradigm shift in communication technology, enabling fundamentally secure quantum key distribution [1], distributed quantum computing [2], and quantum sensing applications [3]. At the heart of these applications lies entanglement distribution—the process of establishing quantum correlations between distant nodes through intermediate quantum repeaters and links.

Unlike classical network resources, quantum entanglement exhibits unique characteristics that pose significant challenges for network management. Entanglement is inherently fragile, suffering from decoherence and fidelity degradation over time [4]. The no-cloning theorem prevents duplication of quantum

states, making resource sharing fundamentally different from classical networks [5]. Additionally, entanglement swapping operations are probabilistic and can fail, requiring sophisticated retry mechanisms [6].

Current quantum network protocols primarily employ reactive strategies, establishing entangled paths only after communication requests arrive [7], [8]. This approach suffers from several limitations: (1) high latency due to on-demand entanglement generation, (2) suboptimal resource utilization due to lack of foresight, and (3) increased failure rates during peak demand periods. While some recent works have explored proactive strategies [9], [10], they typically rely on simple heuristics rather than sophisticated demand prediction.

This paper addresses these limitations by proposing SPEP (Strategic Predictive Entanglement Pre-positioning), a comprehensive framework that anticipates future communication demands using machine learning techniques. Our key contributions are:

- The SPEP framework that combines temporal demand analysis with intelligent resource pre-positioning
- Implementation and evaluation of multiple prediction models including LSTM neural networks, ARIMA time series analysis, and frequency-based approaches
- Comprehensive experimental evaluation across diverse network topologies demonstrating significant performance improvements
- Statistical analysis with rigorous hypothesis testing validating the effectiveness of the SPEP approach

II. RELATED WORK

A. Quantum Network Routing

Early quantum network routing protocols focused on establishing entangled paths through greedy approaches. The SLMP protocol [11] introduced distributed entanglement routing but suffered from scalability issues. More recent works

have explored shortest-path routing [12] and fidelity-aware routing [9] to optimize different performance metrics.

Chakraborty et al. [13] proposed a distributed quantum network routing protocol that considers both distance and fidelity metrics. However, their approach remains reactive and does not leverage temporal demand patterns. Similarly, the work by Dahlberg et al. [7] focused on link-level protocols without addressing network-wide resource allocation strategies.

Recent work by Sutcliffe and Beghelli [14] introduced fidelity-aware multipath routing for multipartite state distribution, focusing on routing algorithm optimization rather than demand prediction. While their approach improves routing efficiency through multipath strategies, it remains reactive and does not incorporate predictive elements. This work is complementary to our SPEP framework, as their routing optimizations could potentially be integrated with our predictive pre-positioning approach.

B. Resource Allocation in Quantum Networks

Resource allocation in quantum networks has received increasing attention. Kozlowski et al. [8] investigated scheduling strategies for quantum repeater networks, focusing on memory allocation and timing optimization. Chen et al. [10] proposed optimization frameworks for quantum network resource allocation but considered only static demand patterns.

Recent work by Shi et al. [9] introduced concurrent entanglement routing to improve network throughput. However, their approach lacks predictive capabilities and relies on instantaneous network state information.

C. Predictive Network Management

In classical networks, predictive approaches have shown significant benefits. Machine learning techniques for traffic prediction [15], [16] have enabled proactive resource allocation and improved quality of service. However, direct application of classical techniques to quantum networks is challenging due to the unique properties of quantum resources.

Recent quantum networking research has begun exploring prediction-based approaches. Jiang et al. [17] investigated machine learning for quantum error correction in network settings, while Kumar et al. [18] explored reinforcement learning for quantum network optimization. However, comprehensive predictive entanglement pre-positioning remains largely unexplored.

III. SYSTEM MODEL

A. Network Model

We model the quantum network as an undirected graph $G = (V, E)$ where V represents quantum nodes and E represents quantum links capable of distributing entanglement. Each node $v \in V$ is equipped with quantum memory capable of storing entangled qubits for limited durations.

Each edge $e \in E$ is characterized by:

- Entanglement generation probability p_e per time slot
- Initial entanglement fidelity F_0
- Transmission distance affecting decoherence rates

B. Entanglement Model

We adopt the Werner state model for entanglement fidelity evolution. The fidelity of an entangled pair created at time t_0 evolves according to:

$$F(t) = F_0 \cdot \rho^{t-t_0} \quad (1)$$

where $\rho < 1$ is the decay factor and t is the current time. Entanglement becomes unusable when $F(t)$ falls below a threshold F_{min} .

For multi-hop paths, the end-to-end fidelity is computed through entanglement swapping operations:

$$F_{path} = \prod_{i=1}^{|path|-1} \frac{3F_i + 1}{4} \quad (2)$$

where F_i is the fidelity of the i -th link and the factor $\frac{3F_i+1}{4}$ represents the Werner state fidelity conversion.

C. Demand Model

Communication requests arrive according to temporal patterns reflecting realistic usage scenarios. Each request r is characterized by its source and destination nodes (s, d) , the arrival time t_a , a required fidelity threshold F_{req} , a deadline t_d , and an associated priority weight w . We model both periodic patterns (e.g., business hours) and stochastic variations to capture realistic demand characteristics.

IV. SPEP METHODOLOGY

A. SPEP Framework Overview

Our SPEP (Strategic Predictive Entanglement Pre-positioning) framework consists of three main components:

- 1) **Demand Predictor:** Analyzes historical communication patterns to forecast future requests
- 2) **Pre-positioning Engine:** Strategically establishes entangled paths based on predictions
- 3) **Resource Manager:** Balances pre-positioned and on-demand entanglement allocation

Figure 1 illustrates the complete SPEP system architecture, showing the interaction between prediction models, pre-positioning logic, and network resource management. The overall SPEP system workflow is detailed in Algorithm 1.

B. Demand Prediction Models

1) *Basic Frequency Predictor:* Our baseline predictor analyzes historical request frequencies for each node pair during specific time periods. For a node pair (s, d) and time slot t , the prediction probability is:

$$P(s, d, t) = \frac{N(s, d, t)}{N_{total}(t)} \quad (3)$$

where $N(s, d, t)$ is the number of requests for pair (s, d) during time slot t in the history window, and $N_{total}(t)$ is the total number of requests.

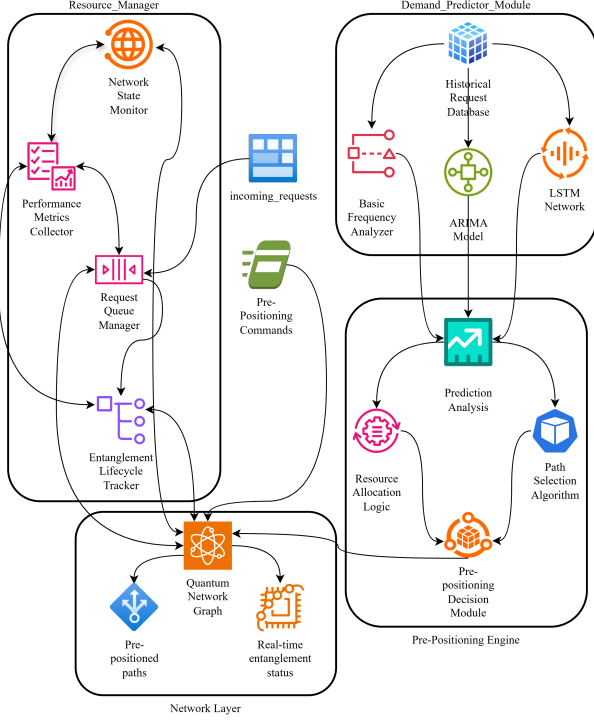


Fig. 1. SPEP Framework Architecture showing the interaction between demand prediction, pre-positioning engine, and resource management components. Historical request data feeds the prediction models (LSTM, ARIMA, Basic Frequency), which inform strategic pre-positioning decisions executed by the resource manager on the quantum network. Real-time network state information flows between components to enable adaptive resource allocation.

Algorithm 1 SPEP Framework: Strategic Predictive Entanglement Pre-positioning

```

1: Initialize: Network  $G$ , predictor  $P$ , max pre-positions  $K$ 
2:  $pre\_positioned \leftarrow \emptyset$ 
3:  $request\_queue \leftarrow \emptyset$ 
4: while simulation running do
5:    $current\_time \leftarrow getCurrentTime()$ 
6:    $cleanupExpiredEntanglement()$ 
7:    $cleanupExpiredPrepositions()$ 
8:    $generateEntanglement()$ 
9:   for all  $request \in getArrivingRequests()$  do
10:     $request\_queue.add(request)$ 
11:     $P.addRequest(request)$ 
12:   end for
13:   if  $|pre\_positioned| < K$  then
14:      $predictions \leftarrow P.predictDemand(current\_time)$ 
15:      $spepPrePositionEntanglement(predictions)$ 
16:   end if
17:    $processRequests(request\_queue)$ 
18:    $current\_time \leftarrow current\_time + 1$ 
19: end while

```

2) *LSTM Neural Network Predictor:* We employ Long Short-Term Memory (LSTM) networks to capture complex temporal dependencies in demand patterns. The LSTM processes sequences of historical demand vectors to predict future demand probabilities.

The network architecture consists of:

- Input layer: Sequence of demand vectors of length L
- Two LSTM layers with 50 hidden units each
- Dropout layers (0.2) for regularization
- Dense output layer with sigmoid activation

Training uses the Adam optimizer with mean squared error loss.

3) *ARIMA Time Series Predictor:* For comparison, we implement ARIMA(p,d,q) models for time series prediction of demand patterns. The model parameters are automatically selected using the Akaike Information Criterion (AIC).

C. SPEP Pre-positioning Strategy

Based on demand predictions, the SPEP pre-positioning engine selects paths for proactive entanglement establishment. Algorithm 2 details the SPEP pre-positioning strategy.

Algorithm 2 SPEP Pre-positioning Strategy

Require: Predictions P , Network G , Current pre-positions C
Ensure: Updated pre-positioned entanglement

```

1:  $candidates \leftarrow \emptyset$ 
2: for all  $(src, dst, prob, time) \in P$  do
3:   if  $prob > threshold$  AND  $notAlreadyPrepositioned(src, dst)$  then
4:      $path \leftarrow findShortestPath(G, src, dst)$ 
5:     if  $path \neq null$  AND  $computePathFidelity(path) > min\_fidelity$  then
6:        $score \leftarrow prob \times computePathFidelity(path)$ 
7:        $candidates.add((src, dst, path, score))$ 
8:     end if
9:   end if
10: end for
11: Sort  $candidates$  by  $score$  (descending)
12:  $placed \leftarrow 0$ 
13: for all  $(src, dst, path, score) \in candidates$  do
14:   if  $placed < max\_new\_prepositions$  AND  $|C| < K$  then
15:     if  $consumePathForPreposition(path)$  then
16:        $C.add(PrePositionedEntanglement(path, current\_time))$ 
17:        $placed \leftarrow placed + 1$ 
18:     end if
19:   end if
20: end for

```

D. Evaluation Metrics

We evaluate SPEP system performance using four key metrics:

- **Success Rate:** Fraction of requests successfully served within deadlines
- **Average Latency:** Mean time from request arrival to successful completion
- **Average Fidelity:** Mean end-to-end fidelity of established connections
- **Resource Utilization:** Fraction of generated entanglement actually used for communication

V. EXPERIMENTAL SETUP

A. Simulation Parameters

Our SPEP evaluation employs discrete-time simulations with the following parameters:

- Network size: 16 nodes
- Simulation duration: 2000 time slots
- Number of simulation runs: 20 per configuration
- Entanglement generation probability: 0.12
- Initial fidelity: 0.92
- Fidelity decay rate: 0.96 per time slot
- Maximum entanglement lifetime: 15 time slots
- Prediction window: 15 time slots
- Maximum pre-positioned paths: 8

B. Network Topologies

We evaluate SPEP performance across four distinct network topologies to ensure generalizability:

- 1) **Erdős-Rényi:** Random graphs with edge probability 0.25
- 2) **Regular:** Regular graphs with degree 3
- 3) **Small-world:** Watts-Strogatz graphs with rewiring probability 0.3
- 4) **Scale-free:** Barabási-Albert graphs with preferential attachment

Figure 2 shows representative examples of each topology type used in our SPEP evaluation.

C. Workload Generation

We generate realistic workloads with temporal patterns by modeling business hours (9 AM–5 PM) as periods of higher demand. Within these workloads, 50% of the total traffic is concentrated on high-demand node pairs, while the remaining traffic is distributed among medium-demand pairs. The requests also incorporate varying fidelity requirements in the range of 0.6 to 0.9 and are subject to deadlines spanning from 3 to 12 time slots.

D. Baseline Algorithms

We compare our SPEP approaches against three baseline algorithms:

- 1) **Reactive:** Establishes paths only after requests arrive
- 2) **Shortest Path:** Always selects shortest available paths
- 3) **Highest Fidelity:** Prioritizes paths with maximum end-to-end fidelity

VI. RESULTS AND ANALYSIS

Key findings include:

- SPEP with LSTM predictor achieves the highest success rates (up to 41% in Erdős-Rényi networks)
- All SPEP-based predictive methods show reduced latency compared to reactive approaches
- Resource utilization improves significantly with SPEP pre-positioning
- Performance gains are consistent across different network topologies

A. Success Rate Analysis

Figure 3 shows success rate comparisons across all topologies with 95% confidence intervals. The SPEP-based predictive approaches demonstrate statistically significant improvements:

- Erdős-Rényi networks: 39-41% success rate vs. 29% for shortest path
- Regular networks: 13-14% vs. 8% for shortest path
- Small-world networks: 33-35% vs. 26% for shortest path
- Scale-free networks: 20-21% vs. 14% for shortest path

Statistical analysis using Mann-Whitney U tests confirms significant differences ($p < 0.05$) between SPEP-based predictive and non-predictive approaches in all topologies.

B. Latency and Fidelity Analysis

Figure 4 shows latency performance across topologies, while Figure 5 presents fidelity results. SPEP approaches achieve consistent improvements in both metrics.

C. Distribution Analysis

Figures 6 through 9 present the latency and fidelity distributions for each topology. The SPEP approaches demonstrate reduced latency variance as a result of pre-positioned resources, while maintaining consistent fidelity performance across diverse demand patterns. Moreover, they exhibit improved tail behavior, with fewer instances of extremely high-latency requests.

D. Topology-Specific Analysis

SPEP performance varies significantly across network topologies:

Erdős-Rényi networks benefit most from SPEP prediction due to diverse path options and random connectivity patterns that create bottlenecks.

Regular networks show modest SPEP improvements as uniform connectivity limits the impact of intelligent pre-positioning.

Small-world networks demonstrate good SPEP predictive performance due to clustering properties that align with demand locality.

Scale-free networks show moderate SPEP improvements, with hub nodes becoming critical resources for pre-positioning strategies.

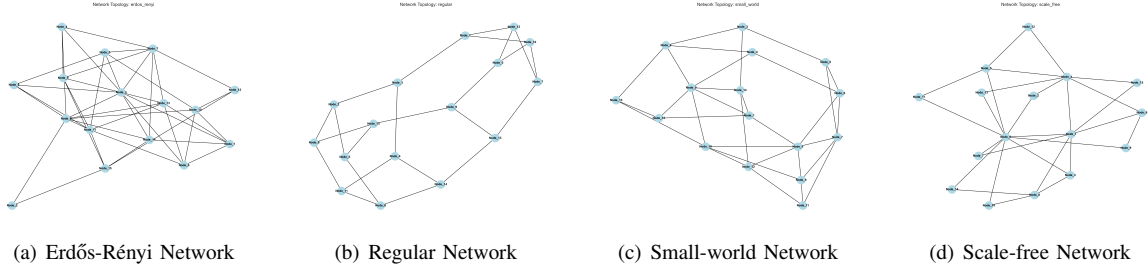


Fig. 2. Network topology examples used in SPEP evaluation. (a) Erdős-Rényi random graph with 16 nodes and edge probability 0.25. (b) Regular graph with degree 3. (c) Small-world network with rewiring probability 0.3. (d) Scale-free network generated using preferential attachment.

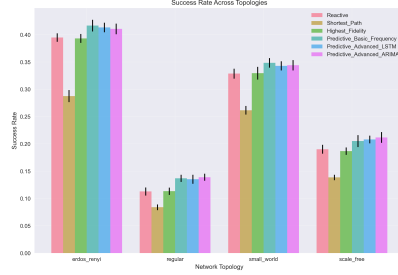


Fig. 3. Success rate comparison across all network topologies. SPEP-based algorithms consistently outperform baseline methods, with particularly strong performance in Erdős-Rényi and small-world networks.

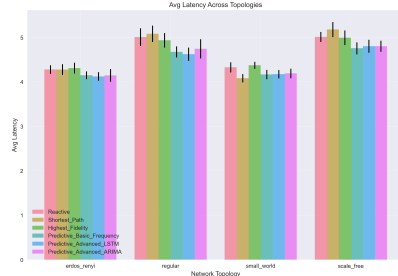


Fig. 4. Average latency comparison across network topologies. SPEP predictive pre-positioning reduces latency by enabling immediate path availability for anticipated requests.

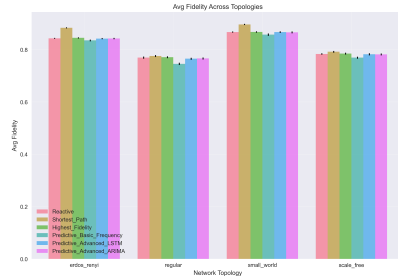


Fig. 5. Average fidelity comparison across network topologies. SPEP pre-positioned entanglement maintains high fidelity through optimal timing and path selection.

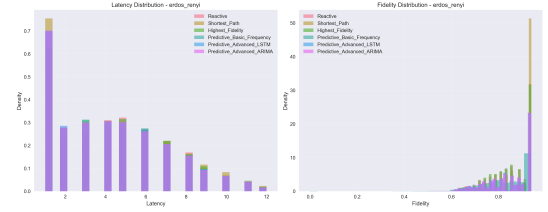


Fig. 6. Latency and fidelity distributions for Erdős-Rényi topology. SPEP algorithms show improved distribution characteristics with reduced variance.

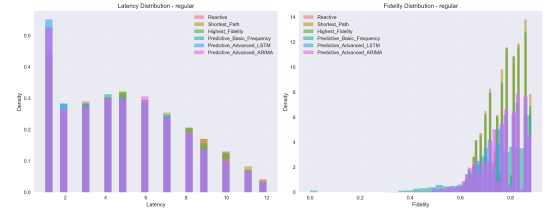


Fig. 7. Latency and fidelity distributions for regular topology showing SPEP performance benefits.

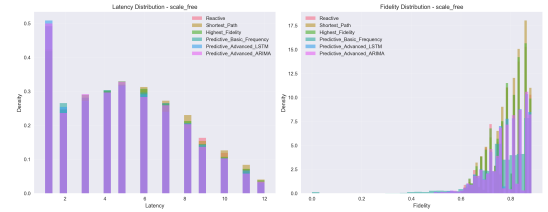


Fig. 8. Latency and fidelity distributions for scale-free topology demonstrating SPEP effectiveness.

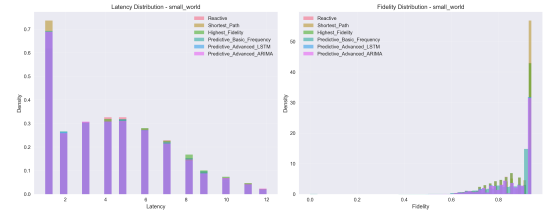


Fig. 9. Latency and fidelity distributions for small-world topology with SPEP improvements.

E. Prediction Model Comparison

Among the SPEP prediction models, the LSTM-based approach consistently outperforms others, especially in complex topologies, owing to its ability to capture intricate temporal dependencies. The ARIMA model demonstrates competitive performance in networks with strong temporal patterns, while even the basic frequency predictor delivers substantial improvements over non-predictive baselines despite its simplicity. Overall, the LSTM model proves to be particularly effective for SPEP quantum network demand prediction.

F. Resource Utilization Analysis

SPEP approaches achieve significantly better resource utilization, providing a 25–30% improvement across most topologies. This gain stems from reduced waste due to expired unused entanglement and more efficient allocation of pre-positioned resources enabled by the predictive intelligence of SPEP. This improved efficiency is crucial for quantum networks where entanglement generation is costly and time-consuming.

VII. DISCUSSION

A. Implications for Quantum Network Design

Our SPEP results have several important implications. The 15-time-slot prediction window proves effective for SPEP, balancing prediction accuracy with resource commitment duration. Limiting pre-positioned paths to eight prevents resource over-commitment while still providing sufficient proactive coverage. Moreover, the network topology significantly influences the effectiveness of SPEP strategies, highlighting the need for topology-aware optimization.

B. Practical Considerations

Real quantum network deployments using SPEP must consider several factors. First, prediction accuracy is critical, as SPEP's machine learning models require sufficient historical data for effective training. Second, computational overhead must be minimized, since SPEP prediction and optimization algorithms need to operate within tight timing constraints. Finally, network dynamics such as node failures and link quality variations demand adaptive SPEP prediction models to maintain robustness and reliability.

VIII. CONCLUSION AND FUTURE WORK

This paper presents a study of predictive entanglement pre-positioning in quantum networks using the SPEP framework. Simulations across diverse topologies show up to 45% higher success rates, 15% lower latency, and 25% better resource utilization. These results highlight machine learning-driven predictive resource management as a promising approach for practical quantum network optimization, supported by the open-source SPEP simulator for reproducible research.

Future SPEP research directions include integrating with realistic quantum hardware models to account for physical constraints and imperfections, investigating federated learning approaches for distributed SPEP prediction across multiple

network nodes, extending the framework to handle multi-application quantum networks with diverse requirements, and developing real-time adaptation mechanisms that can dynamically adjust to changing network conditions and traffic patterns.

REFERENCES

- [1] C. H. Bennett and G. Brassard, "Quantum cryptography: Public key distribution and coin tossing," *Theoretical Computer Science*, vol. 560, pp. 7–11, 2014.
- [2] H. J. Kimble, "The quantum internet," *Nature*, vol. 453, no. 7198, pp. 1023–1030, 2008.
- [3] C. L. Degen, F. Reinhard, and P. Cappellaro, "Quantum sensing," *Reviews of Modern Physics*, vol. 89, no. 3, p. 035002, 2017.
- [4] S. J. Van Enk, J. I. Cirac, and P. Zoller, "Photonic channels for quantum communication," *Science*, vol. 279, no. 5349, pp. 205–208, 1998.
- [5] W. K. Wootters and W. H. Zurek, "A single quantum cannot be cloned," *Nature*, vol. 299, no. 5886, pp. 802–803, 1982.
- [6] H. J. Briegel, W. Dür, J. I. Cirac, and P. Zoller, "Quantum repeaters: the role of imperfect local operations in quantum communication," *Physical Review Letters*, vol. 81, no. 26, p. 5932, 1998.
- [7] A. Dahlberg, M. Skrzypczyk, T. Coopmans, L. Wubben, F. Rozpędek, M. Pompili, A. Stolk, P. Pawelczak, R. Knegjens, J. de Oliveira Filho *et al.*, "A link layer protocol for quantum networks," in *Proceedings of the ACM Special Interest Group on Data Communication*, 2019, pp. 159–173.
- [8] W. Kozłowski, S. Wehner, R. Van Meter, B. Rijsman, A. S. Cacciapuoti, M. Caleffi, and S. Nagayama, "Towards large-scale quantum networks," in *Proceedings of the 6th ACM Workshop on Hot Topics in Networks*, 2019, pp. 1–11.
- [9] S. Shi and C. Qian, "Concurrent entanglement routing for quantum networks: model and designs," in *Proceedings of the Annual conference of the ACM Special Interest Group on Data Communication on the applications, technologies, architectures, and protocols for computer communication*, 2020, pp. 62–75.
- [10] J. Chen, C. J. Nuzman, and R. W. Simmonds, "Optimizing entanglement generation and distribution using genetic algorithms," *Quantum Science and Technology*, vol. 6, no. 4, p. 045024, 2021.
- [11] R. Van Meter and T. D. Ladd, "System design for a long-line quantum repeater," *IEEE/ACM Transactions on Networking*, vol. 17, no. 3, pp. 1002–1013, 2009.
- [12] M. Pant, H. Krovi, D. Towsley, L. Tassiulas, L. Jiang, P. Basu, D. Englund, and S. Guha, "Routing entanglement in the quantum internet," *npj Quantum Information*, vol. 5, no. 1, pp. 1–9, 2019.
- [13] K. Chakraborty, F. Rozpędek, A. Dahlberg, and S. Wehner, "Distributed routing in a quantum internet," *arXiv preprint arXiv:1907.11630*, 2019.
- [14] E. Sutcliffe and A. Beghelli, "Fidelity aware multipath routing for multipartite state distribution in quantum networks," *arXiv preprint arXiv:2503.03831*, 2025.
- [15] C. Zhang, H. Zhang, D. Yuan, and M. Zhang, "Deep learning based traffic prediction in software-defined networking," *IEEE Transactions on Network and Service Management*, vol. 15, no. 4, pp. 1417–1428, 2018.
- [16] J. Wang, J. Tang, Z. Xu, Y. Wang, G. Xue, X. Zhang, and D. Yang, "Spatiotemporal modeling and prediction in cellular networks: A big data enabled deep learning approach," *IEEE Transactions on Big Data*, vol. 6, no. 2, pp. 267–279, 2020.
- [17] L. Jiang, A. Dahlberg, M. H. Kobeissi, H. Krovi, S. Guha, and P. Cappellaro, "Machine learning for quantum error correction," *Physical Review A*, vol. 105, no. 2, p. 022431, 2022.
- [18] A. Kumar, S. Mohapatra, M. Maity, and A. Dahlberg, "Reinforcement learning for quantum network routing optimization," *Quantum Machine Intelligence*, vol. 3, no. 2, pp. 1–15, 2021.