

A GKP qubit-based all-photonic quantum switch

Mohadeseh Azari
School of Computing and Information
University of Pittsburgh
Pittsburgh, PA 15260, USA

Paul Polakos
Cisco Systems
New York, NY 10119, USA

Kaushik P. Seshadreesan
School of Computing and Information
University of Pittsburgh
Pittsburgh, PA 15260, USA

Abstract—We propose and analyze a quantum switch for GKP-qubit-based all-photonic entanglement distribution networks. Its design is compatible with a recently studied design of a GKP-qubit-based all-photonic quantum repeater that achieves high end-to-end entanglement rates despite realistic finite squeezing in the GKP-qubit preparation and homodyne detection inefficiencies. Our main objective is to optimize the allocation of a finite number of multiplexed GKP-qubit-based entanglement resources among different, arbitrary distance client-pair connections enabled by the switch. We achieve this by overcoming the limitations of previous studies by optimizing the bipartite entanglement generation between clients of a switch (or repeater) node even when they are not equally spaced from the switch. We then maximize the switch's total throughput while ensuring the rates are distributed fairly among all client-pair connections. To better illustrate our result, we analyze an exemplary datacenter network where each user aims to connect to the datacenter alone. Together with the quantum repeater, the proposed quantum switch provides a way to realize entanglement distribution-based quantum networks of arbitrary topology.

I. INTRODUCTION

Quantum networking constitutes one of the main pillars of the quantum information revolution that is currently underway [1]. Large distance-scale quantum networks would enable, e.g., secure delegated quantum computation in the cloud [2], secure multiparty quantum computation-based cryptographic protocols [3], [4], and distributed quantum sensing [5]–[7]. Realizing such quantum networks hinges on the ability to communicate quantum information reliably across large distances at high rates, which requires novel quantum node architectures [8], and network infrastructure consisting of specialized quantum-capable helper nodes, namely, quantum repeaters [9] and quantum switches [10].

Light at optical frequencies forms the uncontested best choice of information carrier for quantum communication. However, myriad ways exist to encode quantum information in light, e.g., over polarization modes, time-bins, spatio-spectro-temporal modes, or the continuous quadrature degrees of freedom of individual modes. The choice of encoding strongly affects the design of quantum networks.

Among the different optical quantum information encodings, the Gottesman-Kitaev-Preskill (GKP) qubit encoding [11] is known to be resilient to photon loss. GKP-qubits nearly achieve the quantum capacity of the thermal-noise lossy bosonic communication channel [12] that models common transmission media such as optical fiber and free-space links. As a result, several architectures and protocols

for quantum repeaters based on optical GKP-qubits have been proposed and studied in the context of both entanglement distribution-based networks [13], [14] and forward error-corrected communication-based networks [15].

Here, we concern ourselves with the optical GKP-qubit-based multiplexed all-photonic repeater proposal in ref. [13] on entanglement distribution-based networks (see Fig. 1 for a generalized instance of the architecture). It involves physical GKP-qubit entangled resource states, where the physical GKP-qubits are used for interfacing between nodes, and the logical GKP-qubits consisting of 7 physical GKP-qubits in the $[[7,1,3]]$ Steane code serve as all-photonic quantum memories—the overall state being an 8-qubit graph state of cube topology (up to Hadamards on 4 out of the 8 GKP-qubits). A chain of equispaced repeaters of this type was shown to support end-to-end entanglement rates as high as 0.7 ebits/mode at total distances as large as 700 km under realistic assumptions for GKP-qubit quality expressed in terms of GKP squeezing [16] and coherent homodyne detector efficiencies. Preparing high-quality GKP-qubit-based graph state resources constitutes the primary challenge for these repeaters involving large overheads in terms of the number of physical GKP-qubits required, which calls for their optimal utilization¹.

This paper proposes and analyzes a quantum switch compatible with the above-mentioned quantum repeaters. A quantum switch here refers to a generalized repeater node generating entanglement among neighboring nodes and carrying a switching capability essential to realizing networks of arbitrary topologies. We focus on a bipartite entanglement switch, i.e., a switch that can “connect”, i.e., facilitate bipartite entanglement distribution, between any two among its two or more *clients* (nodes attached to it), that may be most generally at different distances [17]–[19]. Given that the GKP-qubit-based graph state resources form the most valuable commodities at the nodes, the research question this work is answering is how to optimally allocate resources towards the different *connections* that it can enable such that the sum throughput of the switch is maximized. Our work thus completes the prescription for bipartite entanglement distribution network infrastructure based on the quantum repeater proposal of ref. [13].

The paper is organized as follows. In Sec. II, we begin by describing the simplest instance of the proposed quantum

¹Preparing optical GKP-qubits in the first place is challenging too, although there exist proposals based on Gaussian Boson Sampling involving squeezers, multi-mode interferometers, and photon number resolving detectors [16].

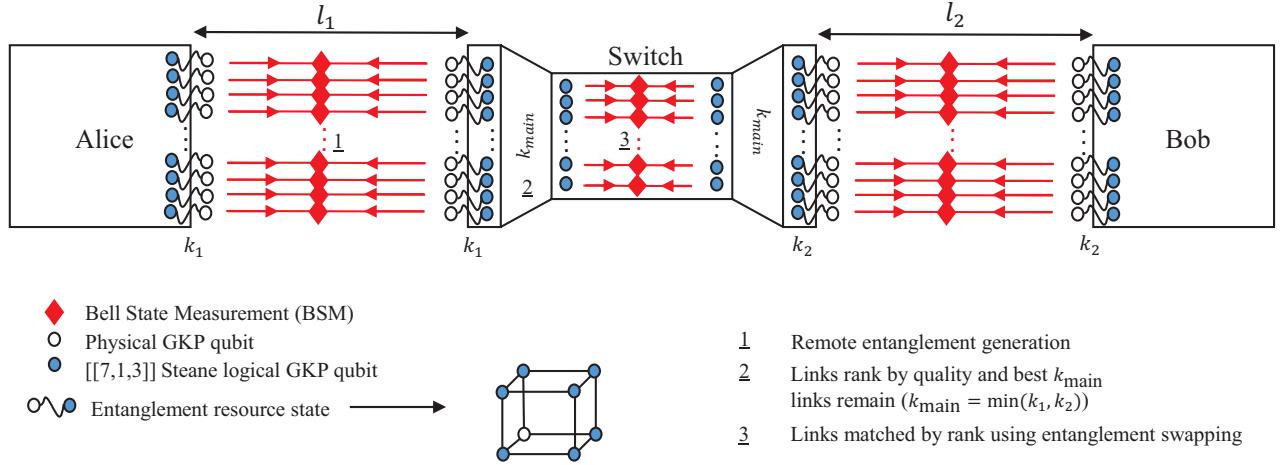


Fig. 1. The proposed multiplexed all-photonic quantum switch based on GKP-encoded qubits and the $[[7,1,3]]$ Steane code in its simplest form consisting of just two clients. The clients are located at distances (l_1, l_2) from the switch. The switch prepares $k_{\text{total}} = k_1(\text{left}) + k_2(\text{right})$ entangled resource states that correspond on the logical level to a Bell pair between a concatenated-coded qubit and a bare GKP-qubit and on the physical level to a cube graph state of eight GKP-qubits, with the clients matching the respective preparations from their end. Remote entanglement generation is performed between the switch and each client by sending the bare physical GKP-qubits (white/empty circle) toward each other for Bell State Measurement (BSM). The elementary entanglement links thus generated are ranked according to their reliability estimated from the GKP-qubit syndromes obtained from the continuous BSM outcomes. This ranking information, as well as logical BSM outcomes, are sent to the switch. The switch chooses the best $k_{\text{main}} = \min(k_1, k_2)$ links from each channel (left and right) to perform entanglement swapping on the concatenated-coded qubits (blue/filled circle) based on that ranking information.

switch, which is all but an improved version of the quantum repeater of ref. [13] including the two clients connection that are most generally at different distances. The section includes a comprehensive mathematical model for the repeater six-state protocol rate and an analysis of its performance. We make key observations about the optimal allocation of resources at the switch for the two clients given a fixed total number of resources and the optimal placement of the switch node given a fixed total distance between clients for achieving the best sum entanglement rates across the repeater. In Sec. III, we describe our general multi-client quantum switch architecture and protocol and elucidate its performance for the instance of datacenter networks. By employing results from Sec. II, we determine the optimum allocation of GKP-qubit-based graph state resources for the different connections that yield the highest total switch rate (sum throughput across the switch). We do so considering fairness between the different entanglement connections enabled by the switch. We conclude with a general set of guidelines for the quantum switch in Sec. IV.

II. QUANTUM SWITCH WITH TWO ASYMMETRICALLY-LOCATED CLIENTS

Consider the simplest instance of the proposed quantum switch as depicted in Fig. 1. The architecture involves two clients that are, in general, located at two different distances (l_1, l_2) and with different numbers of resource states (k_1, k_2) allocated towards each of the clients by the switch to facilitate remote entanglement generation with them, also referred to as *elementary* links. Given a fixed total number of multiplexed resource graph states ($k_{\text{total}} = k_1 + k_2$) at the proposed

switch, the goal is to determine the optimum assignment of resources towards each of the two clients for different values of (l_1, l_2) , where $l_i \in \{0.5, 1, 2, 2.5, 5\}(\text{km})$ for $i \in \{1, 2\}$. The primary objective is to execute the fundamental components (as described in Eq. 6) of the proposed switch and employ the outcomes as a reference point to steer the direction of the key research path.

A. Switch Protocol

To explain the switch protocol for this simple two-client switch, we first describe the repeater protocol of ref. [13]. In a linear chain of equi-spaced repeaters, each pair of neighboring repeater nodes interface with one another via the physical GKP-qubits (referred to as “outer leaf” qubit of the cubic resource graph state) that are transmitted and measured between the nodes by a Bell state measurements (BSM). This results in the generation of logical-logical GKP-qubit-based elementary links in each time step between neighboring nodes. These logical-logical elementary links are retained photonically in local optical fiber spools at the repeater nodes, where the constituent physical qubits (referred as “inner leaf” qubits of the cubic resource graph state) are acted on periodically by quantum error correction, which emulates error-corrected quantum memories. Multiple elementary links of this kind are generated multiplexed in each time step between each pair of neighboring repeaters. A novel feature of the repeater protocol involves an *entanglement ranking*-based link matching strategy for entanglement swapping at the repeater nodes. To elaborate, the elementary links on either side of each repeater node are ranked based on the quality of the generated GKP-qubit

logical-logical entanglement. The quality is inferred from the continuous analog outcome values of the outer leaf GKP-qubit BSMs in terms of

$$P_{\text{no-error}} = (1 - P_p(p_0))(1 - P_q(q_0)), \quad (1)$$

which represents the likelihood of no logical error on the outer leaf qubit during and after the BSM. Here $P_{p/q}$ is the likelihood of facing error, wrongly detecting the GKP-qubit when measuring in the p/q quadrature. Same ranked elementary links along the two clients are then connected by logical-logical GKP-qubit BSMs.

Now, on to the case of the two-client switch, the protocol is modified as follows. Since entanglement-ranking based matching at the switch requires the analog information from the outer leaf BSMs along both clients, given that the clients of the switch are in general at different distances, both sets of inner leaf qubits need to be held in a fiber spool of length $\max(l_1, l_2)$ for entanglement swapping. Further, since the switch may generate different number of multiplexed elementary entanglement (k_1, k_2) with each of the two clients, the switch would then pick the best $k_{\text{main}} = \min(k_1, k_2)$ elementary entanglement links along either clients and performs rank-matched entanglement swapping.

B. End-to-End Entanglement Rate of Switch

We now describe how the performance of the simple two-client switch is evaluated. This involves a generalization of the entanglement distribution rate described in ref. [13] to the case where the elementary links are no longer identical. The total end-to-end rate R_{e2e} depends on a few different quantities, described below. First of these is $Q_{X/Z, \text{outer}, (i)}$, which is a $k_i \times 1$ matrix holding the logical error probability on the outer leaves of the k_i multiplexed elementary entanglement links present between the switch and the i^{th} client. Similarly, we have $Q_{X/Z, \text{inner}, (i)}$ which is a $k_i \times 2$ matrix holding the logical error probability on the inner leaves of the k_i multiplexed elementary links present between the switch and the i^{th} client. Here, whereas the first column represents the logical error when there is no syndrome observed in the Steane code error correction ($s = 0$), the second column represents the logical error when there is a syndrome ($s = 1$).

The total error probability over the j^{th} multiplexed elementary link ($i \in \{1, \dots, k_i\}$) of the switch with the i^{th} client ($i \in \{1, 2\}$) depending on s is given by

$$Q_{X/Z, (i)}(s, j) = Q_{X/Z, \text{inner}, (i)}(s) (1 - Q_{X/Z, \text{outer}, (i)}(j)) + (1 - Q_{X/Z, \text{inner}, (i)}(s)) Q_{X/Z, \text{outer}, (i)}(j). \quad (2)$$

Let $m_{X/Z}$ be how many among the two elementary links measured a non-zero error syndrome on the inner leaves, i.e., $\vec{m}_{X/Z} = (\{m_{X/Z}(i) : i \in \{1, 2\}\})$. Let $\vec{m}_{X/Z}$ be a binary vector of length 2, describing which of the two elementary links had the inner leaf error syndrome during swapping (e.g., $(0, 1)$ means that the inner leaves of the second elementary link, the connection with the second client, had the error

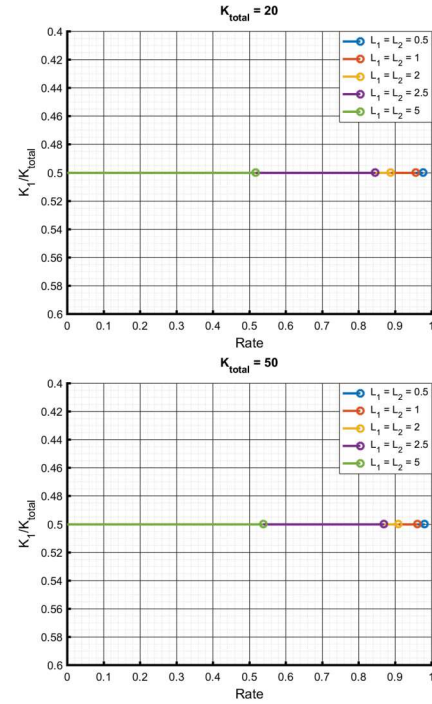


Fig. 2. Rate saturation with the increase in total number of resource states while optimally shared between two clients, i.e., $l_1 = l_2 = l_{\text{total}}/2$ ($l_i \in \{0.5, 1, 2, 2.5, 5\}$ km). For every setting $(l_{\text{total}}, k_{\text{total}})$, the maximum rate per mode $2R_{\text{e2e}}/k_{\text{total}}$ corresponding to the optimum setting (k_1, k_2) is plotted on the X-axis. The optimum allocation, regardless of the distance, is found to be $k_1 = k_2 = k_{\text{total}}/2$, i.e., $k_1/k_{\text{total}} = 0.5$.

syndrome during swapping). For any given $m_{X/Z}$, the set of possible $\vec{m}_{X/Z}$ are given by

$$\|\vec{m}_{X/Z}\|_1 = c_{X/Z} \in \{0, 1, 2\}. \quad (3)$$

Then, the overall end-to-end error probability for the end-to-end link with j^{th} ranking post entanglement swapping at the switch is given by

$$Q_{X/Z, \text{end}}(\vec{m}_{X/Z}, j) = (1/2) \left(1 - \prod_{i=1}^2 (1 - 2Q_{X/Z, (i)}(s=1, j))^{m_{X/Z}(i)} (1 - 2Q_{X/Z, (i)}(s=0, j))^{1-m_{X/Z}(i)} \right). \quad (4)$$

Let $p_{X/Z}(\vec{m}_{X/Z})$ be the probability that the elementary links represented by $\vec{m}_{X/Z}$ did indeed measure a non-zero error syndrome on the inner leaves. It is given by

$$p_{X/Z}(\vec{m}_{X/Z}) = \prod_{i=1}^2 t_{X/Z, (i)}^{m_{X/Z}(i)} (1 - t_{X/Z, (i)})^{1-m_{X/Z}(i)}, \quad (5)$$

where $t_{X/Z, (i)}$ is the probability of an error syndrome ($s = 1$) on the inner leaves.

The total end-to-end entanglement rate of k_{main} multiplexed links based on distilling entanglement separately from different j 's and different $m_{X/Z}$'s is given by

$$R_{\text{e2e}} = \sum_{j=1}^{k_{\text{main}}} \sum_{\vec{m}_X, \vec{m}_Z} p_X(\vec{m}_X) p_Z(\vec{m}_Z) r(Q_{X,\text{end}}(\vec{m}_X, j), Q_{Z,\text{end}}(\vec{m}_Z, j)). \quad (6)$$

where r is the secret-key fraction, a lower bound on distillable entanglement. Together with $Q_{X/Z,\text{inner},(i)} (s = \{0, 1\})$ and $Q_{X/Z,\text{outer},(i)} (j = \{1, \dots, k_i\})$, r is obtained through simulation.

C. Results

We simulated the two-client scenario of the all-optical multiplexed switch architecture with physical GKP-qubits of noise standard deviation of 0.12 (which corresponds to 15dB of GKP squeezing [16]) and detector efficiencies of 0.99, for three values of total resource state $\{k_{\text{total}} = 10, 20, 50\}$. Examining diverse client characteristics has demonstrated that a Gottesman-Kitaev-Preskill (GKP) based architecture is most effective when the system is symmetrical. Any form of asymmetry can hinder the full advantage of the expensive GKP-qubit resource states.

Figure 2 shows that there is an optimum value for the total number of multiplexed links (k_{total}) above which value, increasing the number of resource states does not result in higher end-to-end rate. For each set of (l_1, l_2) , we calculated R_{e2e} for all the possibilities of (k_1, k_2) such that $k_1 + k_2 = k_{\text{total}}$, and ended up with a vector of $\vec{R}_{\text{e2e}}(l_1, l_2)$, where each row contains the end-to-end rate corresponding to a (k_1, k_2) , $k_1 + k_2 = k_{\text{total}}$. In Fig. 3, for each (l_1, l_2) and for a given k_{total} , we plot the configuration that maximizes the rate, i.e.,

$$\begin{aligned} &\text{Maximize } R_{\text{e2e}}(k_1, k_2; l_1, l_2) \\ &\text{s.t. } k_1 + k_2 = k_{\text{total}}, \\ &\quad k_1, k_2 \in \mathbb{Z}^+. \end{aligned} \quad (7)$$

As you can see in the figure, the maximum rate belongs to symmetric distribution, meaning $\text{opt}(k_1, k_2)$ corresponds to $k_1 = k_2 = k_{\text{total}}/2$. If we could show the end-to-end rate of each point through a third dimension, the one where $\{l_1 = l_2, k_1 = k_2\}$ would be the summit (indicated in red diamond in Fig. 3).

Let's study a case where Alice has a distance of $l_1 = 0.5(\text{km})$ whereas Bob has a larger distance of $l_2 = 5(\text{km})$. The effect of increasing the number of multiplexed links of the second client is that the error likelihood of its best link would decrease, meaning Bob will have more "good" links in hand than before. Since Alice is already close to the switch, it guarantees having a small error likelihood for its best link. So by choosing $k_1 < k_2$, $k_1 + k_2 = k_{\text{total}}$, we are improving Bob's outer leaves error. On the other hand, when $k_1 \neq k_2$ it means the switch is throwing away $\{\max(k_1, k_2) - \min(k_1, k_2)\}$ number of entangled resource state, or in other words, it is

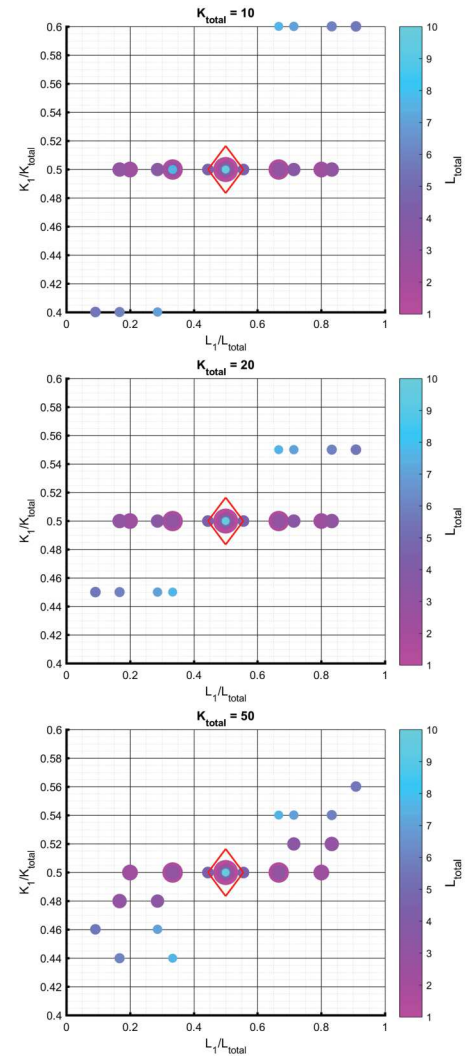


Fig. 3. The general case of a two-client switch where $l_1 + l_2 = l_{\text{total}}$. The total number of entangled resource states (k_{total}) is fixed. There will be an optimum allocation ($\text{opt}(k_1, k_2)$ where $k_1 + k_2 = k_{\text{total}}$) for which the total switch rate is maximum. As you can see in this figure, the optimum allocation is $k_1 = k_2 = k_{\text{total}}$, meaning $k_1/k_{\text{total}} = 0.5$. For every setting $(l_{\text{total}}, k_{\text{total}})$, the optimum number of entangled resource states assigned to the first client over the total number of entangled resource states (k_1/k_{total}) is shown at the Y-axis. The first client distance over the end-to-end distance (l_1/l_{total}) is at the X-axis. This figure only shows the maximum rate over all the (l_{total}) possibilities. The red diamond shows that the maximum rate belongs to $l_1 = l_2 = l_{\text{total}}/2$, $k_1 = k_2 = k_{\text{total}}/2$.

abandoning the costly entangled pairs created in the first step, which is detrimental.

The downside of increasing k_2 outweighs any advantage it may bring because, as mentioned before, resource state generation is the most expensive part of the proposed GKP-qubit-based networking architecture. Thus, employing all the entangled resource states, placing the switch in the middle ($l_1 = l_2$), and allocating links equally ($k_1 = k_2$) achieves the best total end-to-end rate.

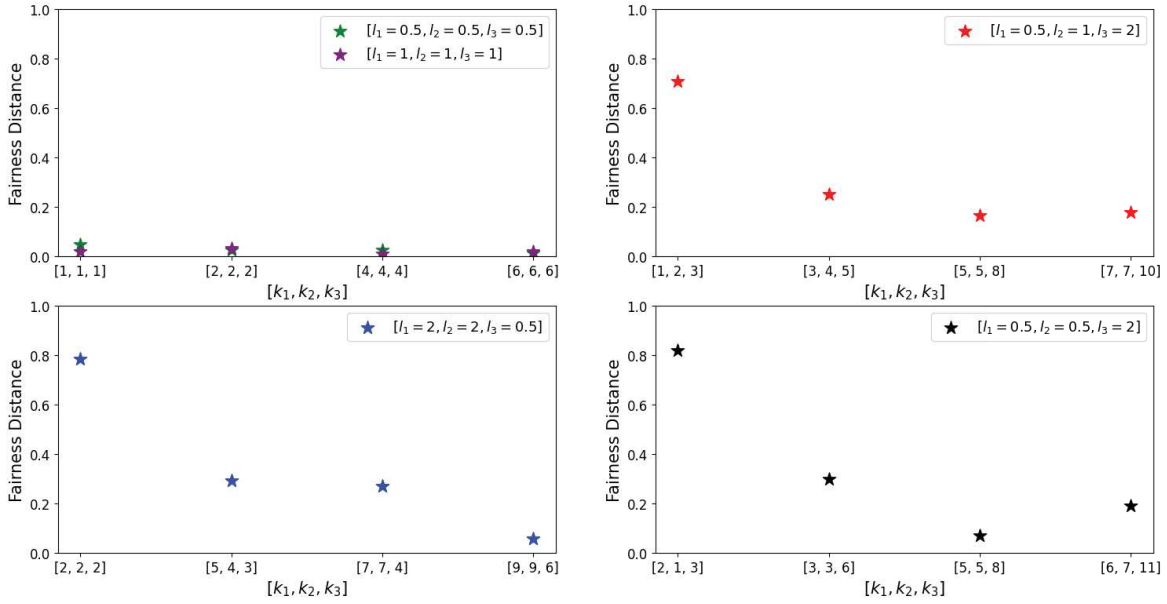


Fig. 4. The best resource allocation adjusted by the switch to ensure the maximum fairness achievable for the given setting $[l_1, l_2, l_3]$ and k_{total} is represented on the x-axis along with the corresponding fairness distance on the y-axis. $[l_1, l_2, l_3]$ are user's distance with $l_4 = 5(\text{km})$ (datacenter distance). With a sufficient k_{total} , we can increase the fairness by assigning more links to the distant user. However, a close user always stands high on the ranking and improves its end-to-end rate by performing entanglement swapping with the best datacenter's links.

III. CONNECTING MULTIPLE CLIENTS

This section focuses on the switch resource allocation problem for a more general multi-client version of the GKP-qubit-based quantum switch. We try to find the optimum allocation of entangled resource states between the multiple client-pair connections that the switch enables. For brevity of analysis, we consider the case where one of the clients in all connections enabled by the switch is always a *datacenter* and the others are *users*. We dismiss all *user-user* connections and solely consider *user-datacenter* connections. The users are also assumed to be closer to the switch than the datacenter ($l_i \ll l_n, i \in \{1, 2, \dots, n-1\}$). The total number of resource states (k_{total}) is assumed to be fixed and is a switch property.

A. Switch Architecture and Protocol

Section II's two-client case showed us that for enabling a client-pair connection with a total resource state of k_{total} , regardless of (l_1, l_2) , the best allocation of resources at the switch that maximizes the rate is always symmetric between the two clients, i.e., $k_1 = k_2$. Given that such is the best allocation of resources within a single connection, the switch protocol treats the different user-datacenter connections that the multi-client switch enables independently and allocates *resources per connection* instead of *resources per user*.

Like in the repeater architecture, the total number of resource states $k_{\text{total}} = \sum_{i=1}^n k_i$ at a switch is assumed to be fixed. All the inner leaves wait for at least $\max(L)$ where $L = (l_1, l_2, \dots, l_n)$ to ensure the switch receives the datacenter outer leaves information. After receiving the ranking information, the switch performs entanglement swapping by connecting the best datacenter's link to the best link on the

user's side, the second-best datacenter's link to the second-best link on the user's side, and so on.

B. Sum Throughput and Fairness of Rates

While optimizing the total sum throughput as above, we also consider a constraint, namely user fairness, which we define as a measure of similarity between rates of different user-datacenter connections. A fair allocation is one with the most semblance between the different user-datacenter rates. Since all the users want to connect to the datacenter, we need to allocate half of the resources anyway to the datacenter ($k_n = k_{\text{total}}/2$). Now the question is: how to allocate the remaining resources ($k_i, i \in \{1, 2, \dots, n-1\}$) to ensure user fairness?

We qualify a resource allocation as fair if

$$d(R_1, R_2, \dots, R_n) < \kappa, \quad (8)$$

where κ is any chosen tolerance, R_i is the rate of the *user_i - datacenter* connection enabled by the switch and $d(R_1, R_2, \dots, R_n)$ is a fairness distance measure that represents the closeness of the rates of the different connections. The rate distance is based on the Euclidean distance between different pairs of rates from the collection of rates $(R_1, R_2, \dots, R_{n-1})$ as

$$d^2(R_1, R_2, \dots, R_n) = \frac{1}{2} \sum_{i,j=1}^n |R_i - R_j|^2. \quad (9)$$

The allocation $(k = [k_1, k_2, \dots, k_{n-1}])$ which gives us the smallest value of fairness distance is chosen as the most fair allocation for the given k_{total} .

C. Results

In our simulation, we set the datacenter distance to be $l_n = 5(\text{km})$ and $l_i \in \{0.5, 1, 2\}(\text{km}), i \in \{1, 2, \dots, n-1\}$. Fig. 4 summarizes the simulations of the switch fairness. Without loss of generality, we considered a simple case with three users and one datacenter at a 5(km) distance from the switch. For each setting $[l_1, l_2, l_3]$ we calculated the switch fairness over four values of k_{total} ($k_{\text{total}} \in \{6, 12, 24, 36, 48\}$). For a given $[l_1, l_2, l_3]$ and k_{total} , the switch decides on an allocation $[k_1, k_2, k_3]$ which is the fairest, in terms of user end-to-end rate ($R_i, i \in \{1, 2, 3\}$). Since we ensure that the datacenter has half of the resource state and the entanglement swap operations at the switch connects links based on their ranking information, the total switch rate stays as high as it is possible to achieve within the feature of entanglement-ranking-based matching for entanglement swapping of the switch protocol.

In Fig. 4, the fairness distance drops as we increase the total number of resource states, meaning the different users' end-to-end rates with the datacenter get close in value. The study on the fairest allocation reveals that for a fair distribution of user-datacenter rates, the user with a further distance from the switch should be allocated more resources than other users. Furthermore, a fair distribution of the user's distance to the switch results in a fair resource allocation. For instance, the $[0.5, 1, 2]$ setting resulted in a fairness distance smaller than the $[0.5, 0.5, 2]$ setting with similar resource allocation. This is because $[0.5, 1, 2]$ is a fairer setting than $[0.5, 0.5, 2]$. The fair allocation of the symmetric settings like $[0.5, 0.5, 0.5]$ or $[1, 1, 1]$ develops an equal end-to-end rate for all the users.

IV. CONCLUSION

To conclude, we presented a GKP-qubit-based multiplexed all-optical quantum switch for entanglement distribution networks. We analyzed the optimal allocation of entanglement resources available at the switch such that the sum throughput of the switch (in other words, the switch rate) is maximized. To maximize the switch rate for a two-client switch, we observed that the switch should be placed in the middle of the connection, and the resources at the switch should be shared equally with the two clients. Symmetric resource allocation always results in a better client-pair rate. In the case of a multi-client switch, we showed that users with a smaller distance from the switch should have fewer resource states to maximize the switch rate while having a fair distribution of rates for the different clients to connect to the datacenter. With a saturation point, the fairness increases as we increase k_{total} . To achieve fairness within any chosen tolerance κ (Eq. 8), a fair setting needs smaller k_{total} .

In an extended version of this work [20], we present a multi-client switch's optimum location and resource allocation concerning total switch throughput and client fairness. Since the first two or three end-to-end links are the ones that mainly drive the end-to-end rate, the switch should be placed where the best links have the smallest outer leaves' errors. If the

clients' locations are fixed, the switch tends to give more resources to the best client to boost its total rate.

ACKNOWLEDGMENT

KPS acknowledges support from NSF grant 2204985 and Cisco Systems. MA and KPS thank Eneet Kaur for helpful discussion.

REFERENCES

- [1] J. P. Dowling and G. J. Milburn, "Quantum technology: the second quantum revolution," *Philos. Trans. A Math. Phys. Eng. Sci.*, vol. 361, no. 1809, pp. 1655–1674, Aug. 2003.
- [2] J. F. Fitzsimons, "Private quantum computation: an introduction to blind quantum computing and related protocols," *npj Quantum Information*, vol. 3, no. 1, pp. 1–11, Jun. 2017.
- [3] V. Lipinska, J. Ribeiro, and S. Wehner, "Secure multiparty quantum computation with few qubits," *Phys. Rev. A*, vol. 102, no. 2, p. 022405, Aug. 2020.
- [4] Y. Dulek, A. B. Grilo, S. Jeffery, C. Majenz, and C. Schaffner, "Secure multi-party quantum computation with a dishonest majority," in *Advances in Cryptology – EUROCRYPT 2020*. Springer International Publishing, 2020, pp. 729–758.
- [5] B. K. Malia, Y. Wu, J. Martínez-Rincón, and M. A. Kasevich, "Distributed quantum sensing with mode-entangled spin-squeezed atomic states," *Nature*, vol. 612, no. 7941, pp. 661–665, Dec. 2022.
- [6] Z. Zhang and Q. Zhuang, "Distributed quantum sensing," *Quantum Sci. Technol.*, vol. 6, no. 4, p. 043001, Jul. 2021.
- [7] X. Guo, C. R. Breum, J. Borregaard, S. Izumi, M. V. Larsen, T. Gehring, M. Christandl, J. S. Neergaard-Nielsen, and U. L. Andersen, "Distributed quantum sensing in a continuous-variable entangled network," *Nat. Phys.*, vol. 16, no. 3, pp. 281–284, Dec. 2019.
- [8] G. Vardoyan, M. Skrzypczyk, and S. Wehner, "On the quantum performance evaluation of two distributed quantum architectures," *SIGMETRICS Perform. Eval. Rev.*, vol. 49, no. 3, pp. 30–31, Mar. 2022.
- [9] W. J. Munro, K. Azuma, K. Tamaki, and K. Nemoto, "Inside quantum repeaters," *IEEE J. Sel. Top. Quantum Electron.*, vol. 21, no. 3, pp. 78–90, May 2015.
- [10] Y. Lee, E. Bersin, A. Dahlberg, S. Wehner, and D. Englund, "A quantum router architecture for high-fidelity entanglement flows in quantum networks," *npj Quantum Information*, vol. 8, no. 1, p. 75, 2022.
- [11] D. Gottesman, A. Kitaev, and J. Preskill, "Encoding a qubit in an oscillator," *Phys. Rev. A*, vol. 64, no. 1, p. 012310, Jun. 2001.
- [12] K. Noh, V. V. Albert, and L. Jiang, "Quantum capacity bounds of gaussian thermal loss channels and achievable rates with Gottesman-Kitaev-Preskill codes," pp. 2563–2582, 2019.
- [13] F. Rozpedek, K. P. Seshadreesan, P. Polakos, L. Jiang, and S. Guha, "All-photonics gottesman-kitaev-preskill-qubit repeater using analog-information-assisted multiplexed entanglement ranking," *Phys. Rev. Res.*, vol. 5, p. 043056, Oct 2023.
- [14] K. Fukui, R. N. Alexander, and P. van Loock, "All-optical long-distance quantum communication with Gottesman-Kitaev-Preskill qubits," *Phys. Rev. Res.*, vol. 3, no. 3, p. 033118, Aug. 2021.
- [15] F. Rozpedek, K. Noh, Q. Xu, S. Guha, and L. Jiang, "Quantum repeaters based on concatenated bosonic and discrete-variable quantum codes," *npj Quantum Information*, vol. 7, no. 1, pp. 1–12, Jun. 2021.
- [16] I. Tzitrin, J. E. Bourassa, N. C. Menicucci, and K. K. Sabapathy, "Progress towards practical qubit computation using approximate Gottesman-Kitaev-Preskill codes," *Phys. Rev. A*, vol. 101, no. 3, p. 032315, Mar. 2020.
- [17] G. Vardoyan, P. Nain, S. Guha, and D. Towsley, "On the capacity region of bipartite and tripartite entanglement switching," *ACM Transactions on Modeling and Performance Evaluation of Computing Systems*, vol. 8, no. 1–2, pp. 1–18, 2023.
- [18] I. Tillman, T. Vasantam, and K. P. Seshadreesan, "A continuous variable quantum switch," in *2022 IEEE International Conference on Quantum Computing and Engineering (QCE)*, 2022, pp. 365–371.
- [19] I. J. Tillman, A. Rubenok, S. Guha, and K. P. Seshadreesan, "Supporting multiple entanglement flows through a continuous-variable quantum repeater," *Physical Review A*, vol. 106, no. 6, p. 062611, 2022.
- [20] M. Azari, P. Polakos, and K. P. Seshadreesan, "Quantum switches for gottesman-kitaev-preskill qubit-based all-photonics quantum networks," *arXiv preprint arXiv:2402.02721*, 2024.