

This is a repository copy of *Practical Routing and Criticality in Large-Scale Quantum Communication Networks*.

White Rose Research Online URL for this paper:

<https://eprints.whiterose.ac.uk/id/eprint/233147/>

Version: Accepted Version

Preprint:

Harney, Cillian and Pirandola, Stefano orcid.org/0000-0001-6165-5615 (2025) Practical Routing and Criticality in Large-Scale Quantum Communication Networks. [Preprint]

<https://doi.org/10.48550/arXiv.2509.10908>

Reuse

Items deposited in White Rose Research Online are protected by copyright, with all rights reserved unless indicated otherwise. They may be downloaded and/or printed for private study, or other acts as permitted by national copyright laws. The publisher or other rights holders may allow further reproduction and re-use of the full text version. This is indicated by the licence information on the White Rose Research Online record for the item.

Takedown

If you consider content in White Rose Research Online to be in breach of UK law, please notify us by emailing eprints@whiterose.ac.uk including the URL of the record and the reason for the withdrawal request.

Practical Routing and Criticality in Large-Scale Quantum Communication Networks

Cillian Harney^{1,2} and Stefano Pirandola^{1,*}

¹*Department of Computer Science, University of York, York YO10 5GH, UK*

²*nodeQ, 71-75 Shelton Street, Covent Garden, London WC2H 9JQ, UK*

The efficacy of a communication network hinges upon both its physical architecture and the protocols that are employed within it. In the context of quantum communications, there exists a fundamental rate-loss tradeoff for point-to-point quantum channels such that the rate for distributing entanglement, secret keys, or quantum states decays exponentially with respect to transmission distance. Quantum networks are the solution to overcome point-to-point limitations, but they simultaneously invite a challenging open question: How should quantum networks be designed to effectively and efficiently guarantee high rates? Now that performance and physical topology are inexorably linked, this question is not easy, but the answer is essential for a future quantum internet to be successful. In this work, we offer crucial insight into this open question for complex optical-fiber quantum networks. Using realistic descriptions of quantum networks via random network models and practical end-to-end routing protocols, we reveal critical phenomena associated with large-scale quantum networks. Our work reveals the weaknesses of applying single-path routing protocols within quantum networks, observing an inability to achieve reliable rates over long distances. Adapting novel algorithms for multi-path routing, we employ an efficient and practical multi-path routing algorithm capable of boosting performance while minimizing costly quantum resources.

I. INTRODUCTION

The deployment of a world-wide quantum network will enable provably-secure communication and cryptography [1–6], and facilitate distributed quantum computing [7–9] on a global scale. A quantum internet will provide new means of information processing, leading to enhancements in a number of technological domains [10–12].

However, in many ways a quantum internet will not be like its classical counterpart [13–17]. Unlike classical information which is robust and easily copied, quantum information is fragile and the laws of quantum mechanics prohibit cloning. As a result, the rate at which one can perform entanglement distribution, quantum key distribution (QKD) or quantum state transfer is fundamentally limited by the medium through which it travels. The optimal transmission rate through optical fiber (the primary conduit for communications) is known exactly as the Pirandola-Laurenza-Ottaviani-Banchi (PLOB) bound [19] which decays exponentially with respect to transmission distance [18]. Crucially, the infrastructure necessary for reliable quantum communications are not immediately compatible with current classical constructions, motivating numerous contributions in the literature [20–29].

The fundamental limits of *end-to-end* communication in quantum networks are also well understood [21]. Optimal performance in an end-to-end setting is achieved by a network with capacity achieving links operating a flooding protocol; a protocol in which every channel in the network is used to facilitate communication between two end-users. These tools have been explored and expanded

in a number of works, including the development of analytical tools for large-scale networks [30–32] and experimental demonstrations in QKD networks [33]. Zhuang *et al.* carried out invaluable assessments of random, optical-fiber quantum networks [34, 35]. Random network models are capable of capturing complex behaviours that arise within real world architectures [36, 37]. Hence, understanding how quantum communications can be performed effectively across such models provides insight into desirable properties for the quantum internet.

Nonetheless, glaring gaps exist in our understanding of realistic performance and resource demands. The most expensive resource in a quantum internet will be the number of quantum repeaters needed in a physical network area to guarantee effective communications, i.e. the network nodal density. While classical repeaters are cheap and easy to deploy, quantum repeaters are costly and should not be wasted as non-user nodes. Recent studies have been able to identify important nodal density conditions for reliable end-to-end performance on quantum networks [30, 34, 38]. Of these investigations, Ref. [38] principally focuses on network connectivity properties and a specific single-photon transmission protocol through pure-loss channels, while Refs. [30, 34] evaluate network capacities with a protocol agnostic approach, using pure-loss channels and optimal flooding protocols.

There is room for progress in each of these works. Network connectivity investigations are useful, but are not sufficient to benchmark end-to-end performance. Beyond pure-loss models, optical fiber is most accurately described using thermal-loss channels that account for environmental and experimental thermal noise. Finally, the employment of flooding protocols on a large-scale is highly impractical. It is not realistic to assume the use of *every* edge in a potentially global network to fa-

* Corresponding author; stefano.pirandola@york.ac.uk

cilitate communication between a single end-user pair. Some of these assumptions are effective and completely suitable for network capacity assessments, they are over-optimistic in the context of revealing stricter, practical insight for quantum networking.

In this work, we tighten the assessment of critical resource requirements in realistic quantum networks with two crucial improvements. Firstly, we consider random network architectures composed of realistic link-layers; using point-to-point channels which account for thermal noise, experimental imperfections and practical point-to-point protocols. This immediately improves our characterisation of realistic quantum networks, and reveals the impact of point-to-point limitations on network connectivity models.

Secondly, we explore the efficacy of practical end-to-end routing in quantum networks to identify critical resource requirements of quantum networks. Moving away from impractical flooding schemes, we study the feasibility of single-path routing and resource efficient multi-path routing, adopting recent results in the efficient generation of end-to-end multi-paths [39]. Our results suggest that single-path routing is an inefficient strategy for reliable rates on large-scale quantum networks. Instead, there exist efficient multi-path routing protocols which can efficiently achieve close to optimal rates without consuming the entire network. Furthermore, different network architectures observe a suitability to quantum multi-path routing which others do not. All of these results identify important design criteria for future, large-scale quantum networks while motivating the further development of practical multi-path protocols.

II. QUANTUM NETWORKS

A. Preliminaries

A quantum communication network can be described as a collection of nodes $P = \{\mathbf{x}_i\}$ which are interconnected by a collection of edges $E = \{(\mathbf{x}_i, \mathbf{x}_j)\}_{i,j}$, resulting in a finite, undirected graph $\mathcal{N} = (P, E)$. A network node $\mathbf{x} \in P$ represents a station which contains a register of quantum systems and the capability to transmit these systems to other nodes in the network. This station may act as a repeater node, or a potential user of the network. A pair of nodes $\mathbf{x}$ and $\mathbf{y}$ are connected by an undirected edge $(\mathbf{x}, \mathbf{y}) \in E$ if there exists a quantum channel $\mathcal{E}_{\mathbf{x}\mathbf{y}}$ between them; permitting the exchange of quantum systems between them.

A quantum network can be used to facilitate many different communication configurations which fall into two main categories: End-to-end and multi-end communications. In an end-to-end network protocol, two remote users $\mathbf{a}, \mathbf{b} \in P$ aim to establish quantum communication between one another, and all other nodes in the network can be used to enable this goal. Multi-end communication configurations form a much broader class

with tasks ranging from multiple-unicast (many pairs of end-to-end protocols simultaneously executed on the same network) to multi-casting (senders communicating unique messages with many receivers) and broadcasting (senders communicating identical messages with many receivers).

In this work, we focus on end-to-end strategies. The ability of a quantum network to achieve strong end-to-end rates serves as a primitive for all other communication configurations. If strong end-to-end rates cannot be achieved, this promises difficulties for more complex multi-end scenarios in which network competition will only serve to reduce performance.

B. Network Link Layers

For an arbitrary quantum network $\mathcal{N} = (P, E)$ we can define a rate distribution $\mathcal{K}$ as the collection of point-to-point rates $\mathcal{K} := \{K_{\mathbf{x}\mathbf{y}}\}_{(\mathbf{x}, \mathbf{y}) \in E}$, where $K_{\mathbf{x}\mathbf{y}}$ denotes the rate at which a network edge is able to operate. The rate distribution of any network depends on the end-to-end task at hand (key distribution, entanglement distribution, quantum state transfer), the channel quality and the point-to-point protocols being used. Here, we focus on realistic description of the link layer via bosonic thermal-loss channels and consider two distinct descriptions of their rates: (i) Each link operates at its two-way assisted quantum and private capacities, delivering point-to-point protocol-agnostic insight to the limits of the network, and: (ii) The rate of each link is described by an asymptotic secret key-rate according to practical QKD protocols, providing insight into the properties of future QKD networks.

We begin by considering quantum channel capacities. Bosonic quantum communications performed over optical-fiber is best described using a Gaussian thermal-loss channel $\mathcal{E}_{\eta, \bar{n}}$ of transmissivity $\eta \in (0, 1)$ and output thermal noise of $\bar{n}$ photons. This channel can be described by the action of a beam-splitter of transmissivity η which mixes the input mode with an environmental thermal mode with mean photon number $\bar{n}_{\text{env}} := \bar{n}/(1 - \eta)$ [40]. This transforms the input quadratures of a single-mode input Gaussian state $\hat{\mathbf{x}} = (\hat{q}, \hat{p})^T$ according to $\hat{\mathbf{x}} \rightarrow \sqrt{\eta}\hat{\mathbf{x}} + \sqrt{1 - \eta}\hat{\mathbf{x}}_{\text{env}}$, where $\hat{\mathbf{x}}_{\text{env}}$ is the quadrature operator of the thermal environmental mode. When $\bar{n} = 0$, the environmental mode is in the vacuum state, and this becomes a pure-loss channel.

Unfortunately, the exact capacity of bosonic thermal-loss channels is not known but can instead be tightly bounded. Defining the channel transmissivity $\eta_{\text{ch}} := 10^{-\gamma d}$ where $\gamma \approx 0.02$ (corresponding to 0.2 dB/km) is the fiber loss rate and d is the channel length (km), we can write the bounds $\mathcal{T}_{\eta, \bar{n}}^l \leq \mathcal{C}(\mathcal{E}_{\eta, \bar{n}}) \leq \mathcal{T}_{\eta, \bar{n}}^u$, where we define the bounding functions,

$$\mathcal{T}_{\eta, \bar{n}}^l := -\log_2(1 - \eta) - h(\bar{n}_{\text{env}}), \quad (1)$$

$$\mathcal{T}_{\eta, \bar{n}}^u := \mathcal{T}_{\eta, \bar{n}}^l - \bar{n}_{\text{env}} \log_2(\eta), \quad (2)$$

and $h(x) := (x+1)\log_2(x+1) - x\log_2(x)$ is an entropic function.

Hence, the optimal rate of any end-to-end protocol on a bosonic thermal-loss network can always be bounded by modelling its rate distributions using the single-edge capacity bounds. That is, given a network $\mathcal{N} = (P, E)$, we may investigate the performance of an end-to-end protocol on this network subject to the rate distributions,

$$\mathcal{K}_l = \{\mathcal{T}_{\eta_{xy}, \bar{n}_{xy}}^l\}_{(x,y) \in E}, \quad (3)$$

$$\mathcal{K}_u = \{\mathcal{T}_{\eta_{xy}, \bar{n}_{xy}}^u\}_{(x,y) \in E}. \quad (4)$$

Using this knowledge, we can bound the optimal performance of end-to-end protocols on bosonic thermal-loss networks. A similar approach can of course be applied to networks composed of other channels with unknown channel capacities.

Unfortunately, practical, near-term performance levels for quantum communications remain some way off of optimal rates described by channel capacities. It is therefore beneficial to investigate the properties of quantum networks subject to realistic protocols and technologies that are achievable today. In this regard, QKD stand as the most deployable quantum communication task to date; yet, realistic resource assessments of QKD networks remain largely unknown. Hence, we also initiate an assessment of large-scale trusted-node QKD networks. By treating each network node as a trusted user station, we are able to investigate large-scale network constructions in which the single-edge rates are captured by well known point-to-point secret-key rates. In this way, we are able to exploit practical key-rates associated with QKD protocols and assess the capabilities/resource requirements of QKD networks. To see the precise key-rate expressions, we point the reader to Appendix A.

C. End-to-End Routing and Rates

An end-to-end protocol is characterised by its routing strategy, the protocol which dictates the path (or paths) that logical quantum communication [43] follows throughout the network in order to communicate between end-users. An end-to-end route between a pair of remote end-users $\mathbf{a}$, $\mathbf{b}$ can be defined as a collection of network edges which connect one end-user to the other. We can define a single-path route as a set

$$\omega := \{(\mathbf{a}, \mathbf{x}_1), (\mathbf{x}_1, \mathbf{x}_2), \dots, (\mathbf{x}_N, \mathbf{b})\}, \quad (5)$$

and define a route as free of cycles (without loss of generality). Any route ω refers to an associated collection of quantum channels $\{\mathcal{E}_{xy}\}_{(x,y) \in \omega} = \{\mathcal{E}_{\mathbf{a}\mathbf{x}_1}, \mathcal{E}_{\mathbf{x}_1\mathbf{x}_2}, \dots, \mathcal{E}_{\mathbf{x}_N\mathbf{b}}\}$, through which quantum systems must be exchanged to establish end-to-end quantum communications. In quantum networks, a routing strategy must be chosen to dictate how network interactions are followed. There exist two main classes of end-to-end routing strategy: *Single-path* and *multi-path* routing.

Single-path routing describes a network protocol, $\mathcal{P}_{\text{sp}}$, in which quantum systems are exchanged from node-to-node throughout the network in a sequential manner. This forges a unique path of interactions through the network, and continues until quantum communication has been established between the end-users. This is the standard strategy used for classical networking and is extremely effective thanks to the robustness of classical information.

The vulnerability of quantum networks to decoherence means that it is extremely valuable to explore more general multi-path protocols, $\mathcal{P}_{\text{mp}}$ in order to enhance performance. Network nodes may exchange multiple quantum systems with many receiver nodes, repeating until communication is established between the end-users. End-users may explore a variety of end-to-end routes simultaneously in a way which enhances their performance.

Let us give a generalised description of an end-to-end protocol. Consider a quantum network $\mathcal{N} = (P, E)$, a routing protocol $\mathcal{P}$ and an end-user pair Alice $\mathbf{a}$ and Bob $\mathbf{b}$ which we collect into a single end-user quantity for ease of notation, $\mathbf{i} := \{\mathbf{a}, \mathbf{b}\}$ ($\mathbf{a}, \mathbf{b} \in P$). Alice prepares a generally multipartite quantum state in her register and performs a point-to-multipoint transmission, exchanging quantum systems with all the nodes in her neighbourhood $\mathbf{a} \rightarrow n_{\mathbf{a}} := \{\mathbf{x} \in P \mid (\mathbf{a}, \mathbf{x}) \in E\}$, where $n_{\mathbf{x}}$ denotes the neighbourhood of a node $\mathbf{x}$. We assume that the use of any edge is probabilistic, such that any node $\mathbf{x}$ only exchanges systems with a node $\mathbf{y}$ according to a *forwarding probability*, $q_{xy} \in [0, 1]$. After this point-to-multipoint transmission, each node $\mathbf{x} \in n_{\mathbf{a}}$ will do the same thing with each of their neighbours and perform system exchanges along any available edge which has not yet been used (according to their own forwarding probability). This process of multipoint exchanges will continue throughout the network until eventually a final multipoint-point exchange will occur with Bob's node and his neighbourhood $n_{\mathbf{b}} \rightarrow \mathbf{b}$. At this point, communication has been established between the end-users.

Over the course of many transmissions, every edge $(\mathbf{x}, \mathbf{y}) \in E$ in the network will have been utilised q_{xy} -times on average. In this way, we have described a generalised version of the standard *flooding protocol* $\mathcal{P}_{\text{f}}$ [21]. Instead of using every edge once in order to establish end-to-end communication, i.e. $q_{xy} = 1$ for all $(\mathbf{x}, \mathbf{y}) \in E$, each edge is used $q_{xy} \leq 1$ times on average.

In large networks, the use of standard flooding is not practical, since this requires the use of potentially huge number of edges for a single-pair of uses. The distribution of forwarding probabilities $\{q_{xy}\}_{(x,y) \in E}$ is completely defined by the chosen routing protocol. It can be readily modified to exclusively perform end-to-end communication along a selection of $M \geq 1$ paths which we collect into the route set, $\Omega_{\mathcal{P}} := \{\omega_1, \omega_2, \dots, \omega_M\}$. It is useful to denote an analogous routing edge set $E_{\mathcal{P}}$ which unpacks

the route set into all its unique edges,

$$E_{\mathcal{P}} := \omega_1 \cup \dots \cup \omega_M = \bigcup_{i=1}^M \omega_i. \quad (6)$$

More precisely, a protocol $\mathcal{P}$ which deterministically uses a finite set of routes $\Omega_{\mathcal{P}}$ generates a forwarding probability distribution of the form

$$q_{\mathbf{x}\mathbf{y}} = \begin{cases} 1 & \text{iff } (\mathbf{x}, \mathbf{y}) \in E_{\mathcal{P}}, \\ 0 & \text{otherwise,} \end{cases} \quad (7)$$

for all $(\mathbf{x}, \mathbf{y}) \in E$. It is clear that if we reduce this set of potential paths to only one route $\Omega_{\mathcal{P}} = \{\omega_1\}$ then we reclaim a single-path protocol, $\mathcal{P}_{\text{sp}}$. Otherwise, we engage in a multi-path protocol, $\mathcal{P}_{\text{mp}}$.

The end-to-end rate associated with this general form of protocol can be computed by solving the max-flow min-cut theorem [21]. That is, an end-to-end rate is found by determining the set of edges in the network which simultaneously disconnect the end-user pair and minimize the sum of all their single-edge rates. We call this partitioning a *network cut* C which generates the set of edges which creates the partition, or cut-set $\tilde{C}$ (see Appendix B). Given a single-edge rate $K_{\mathbf{x}\mathbf{y}}$ and its forwarding probability $q_{\mathbf{x}\mathbf{y}}$, its effective rate is given by $q_{\mathbf{x}\mathbf{y}}K_{\mathbf{x}\mathbf{y}}$. Hence, the end-to-end rate between i on the network $\mathcal{N}$ using $\mathcal{P}$ can be computed by

$$K(i, \mathcal{N}|\mathcal{P}) := \min_C \sum_{(\mathbf{x}, \mathbf{y}) \in \tilde{C}} q_{\mathbf{x}\mathbf{y}} K_{\mathbf{x}\mathbf{y}}. \quad (8)$$

III. RANDOM QUANTUM NETWORKS

A. Link Quality and Edge Pruning

The study of random network models is a rich and wide-spanning field, within which many key tools have been developed. Of course, random network models cannot capture *all* of the features of a future quantum internet; these features will emerge as the technology is developed and deployed. Nonetheless, relevant random networks are able to capture and predict important behaviours of realistic complex networks. In this work, we consider the classes of Waxman and scale-free networks with two important considerations: Realistic link-layer descriptions and practical end-to-end protocols.

To address the weak rates achieved by quantum links over long distances, we impose a modification to the standard random network models. Network edges which possess a point-to-point rate below some threshold value ε are *pruned* and removed from the network edge set after graph generation. Under this modification networks can only be considered completely connected if they are *competently connected* by edges with $K_{\mathbf{x}\mathbf{y}} \geq \varepsilon$. Given that desirable clock rates C for quantum communications are

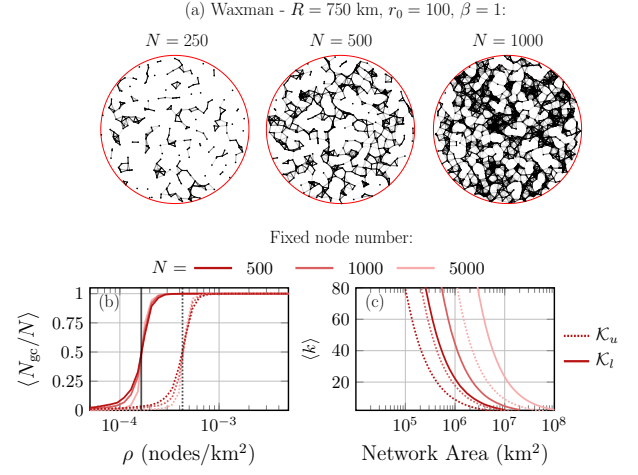


Figure 1. Connectivity properties of bosonic thermal-loss Waxman networks. Panel (a) displays random networks generated under the parameters listed above and using single-edge capacity upper-bounds. The colour intensity of each edge is proportional to its capacity. Panel (b) plots the average main component fraction, in which the critical densities necessary for a connectivity phase transition are identified by the black lines. Panel (c) shows the average degree of networks with a fixed number of nodes (given in legend) with respect to variable network density and area respectively. The legend identifies the fixed node number and indicates whether the network uses upper or lower bounds on the rate distributions.

on the order of GHz, we set $\varepsilon \sim 10^{-12}$ so that valid network edges must guarantee at least $CK_{\mathbf{x}\mathbf{y}} \gtrsim 1$ mbit per second [44]. The notion of pruning goes towards preserving network resources that are otherwise wasted, and offers a more accurate representation of network connectivity. Importantly, it affects our random network models in unique ways, reflecting the potential impact of poor link quality in a realistic quantum internet.

B. Waxman Networks

Networks from the Waxman class $\mathcal{N} \in \mathcal{W}$ are constructed as follows: A number of N nodes are generated in a region of area A , and any pair of nodes $\mathbf{x}, \mathbf{y}$ are connected with a probability that decays exponentially with respect to their point-to-point separation, $r_{\mathbf{x}\mathbf{y}}$. More precisely, a channel $\mathcal{E}_{\mathbf{x}\mathbf{y}}$ between the nodes $\mathbf{x}, \mathbf{y}$ is created with probability

$$p_{\mathbf{x}\mathbf{y}}^{\mathcal{W}} := \beta e^{-\frac{r_{\mathbf{x}\mathbf{y}}}{r_0}}. \quad (9)$$

The parameters β, r_0 are characteristic of the model and influence generation process; $\beta \in (0, 1]$ defines a maximum probability of connection, while $r_0 \in (0, \infty)$ dictates the speed at which the connection probability exponentially decays. Erdős-Rényi networks are a specification of the Waxman model such that $r_0 \rightarrow \infty$, i.e. there is no decay with respect to channel length but the probability of connection is simply $p_{\mathbf{x}\mathbf{y}}^{\text{ER}} := \beta \leq 1$.

Examples of N node networks contained in circular areas of radius R can be seen in Fig. 1(a). In order to analyse these networks, we define N_{gc} as number of nodes in the giant component of the network (largest subset of nodes between which all nodes possess paths to one another). Furthermore, we define the degree of a node k_x as the number of nodes to which it is connected. When N is very small (or R is very large) the network is sparse and poorly connected; the average distance between nodes is large so that links will connect with very low probability under Eq. (9). As a result, the network will be clustered and incompletely connected [45]. This can be seen in Fig. 1(a) for $N = 250$ nodes. Yet, as N is increased (or R is decreased) then the average nodal separation will shrink, increasing the likelihood of connections. Eventually, the nodal density becomes large enough to promise that the number of nodes in the giant component $N_{gc} = N$ and there exists end-to-end paths between all nodes.

This behaviour in Waxman networks is well known, and gives rise to a percolation phase transition. As seen in Fig. 1(b) There exists a critical nodal density ρ_G^* at which the model abruptly transits from poorly connected ($N_{gc} < N/2$) to well connected ($N_{gc} \geq N/2$). The phase transition is illustrated in Fig. 1(b) in the context of bosonic thermal-loss networks, where we plot the average fraction of the network contained in the giant component, $\langle N_{gc}/N \rangle$. One can similarly see in Fig. 1(c) how the average nodal degree $\langle k \rangle$ undergoes a collapse as the network area is expanded, leading to the connectivity transition. It can be seen that the critical density is bounded between $1.6 \times 10^{-4} \lesssim \rho_G^* \lesssim 4.3 \times 10^{-4}$ nodes per km^2 . This is nearly two orders of magnitude larger than that predicted by bosonic pure-loss networks, for which $\rho_G^* \sim 7 \times 10^{-6}$ nodes per km^2 [34, 38]. This emphasises thermal decoherence as a vital consideration since the impact of environmental noise alone can significantly degrade connectivity.

C. Scale-free Networks

Another important random architecture is that of scale-free networks S . A network is scale-free if their degree distribution (probability distribution of nodes in S having degree k) follows a power law, i.e. $p_k \propto k^{-\gamma}$ where γ is some real number characterizing the distribution. Many real world networks (such as the classical internet) are thought to exhibit scale-free properties, however it is rare that a network is precisely scale-free [46]. In this work, we study scale-free networks generated dynamically using Yook's model [47, 48] in which new nodes y are iteratively added to an initially small, n_0 node connected network. Each new node y is attached to a collection of m existing nodes $\{x_i\}_{i=1}^m$ with probability

$$p_{xy}^S \propto k_x^{\sigma_{deg}} / r_{xy}^{\sigma_r}, \quad (10)$$

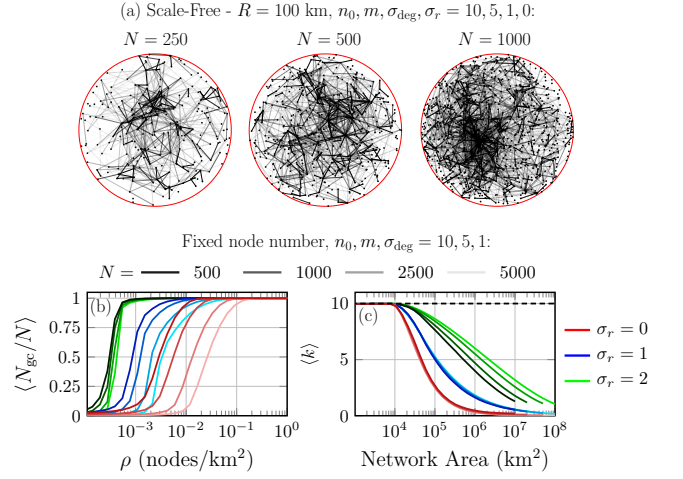


Figure 2. Connectivity properties of bosonic thermal-loss scale-free networks. Panel (a) displays random networks generated under the parameters listed above and using single-edge capacity upper-bounds. The opacity of each edge is proportional to its capacity. Panel (b) plots the average giant component fraction and (c) the average degree of networks with a fixed number of nodes with respect to variable network density and area respectively. The legend identifies the fixed node number and σ exponent.

where $\sigma_{deg}, \sigma_r \in \mathbb{R}_0^+$ are model parameters which controls the influence degree and link length have on connection probability.

Some example networks are illustrated in Fig. 2(a). In general, the connectivity properties of scale-free networks are very different to those of Waxman networks. Fig. 2(b) shows the behaviour of the giant component with respect to fixed node number and nodal density. When there is no connection dependence on link length ($\sigma_r = 0$) there is no critical density. The giant component eventually transits to $\langle N_{gc}/N \rangle = 1$, but less abruptly, and at a higher density than that of Waxman networks. Increasing σ_r , one can see that the network becomes fully connected more quickly and a critical density begins to emerge as we recover the transition shape from Fig. 1(b).

There are also notable differences in the context of average nodal degree, $\langle k \rangle$. The average degree of Waxman networks scales exponentially with respect to nodal density. In the scale-free setting, the relationship between connection probability and nodal degree gives rise to nodal hubs (i.e. nodes to which most others are connected) but also a significant level of sparsity, so that most nodes have a very low degree. This is due to the mechanism of preferential attachment, as nodes which have high degrees are more likely to gain more connections. As seen in Fig. 2(c), at high densities, scale-free networks saturates to a maximum value $\langle k \rangle \rightarrow 2m$.

However, at lower network densities this is not the case and the average degree undergoes a collapse. This can be understood as follows: When new nodes are added, they attempt to connect with m other existing network

nodes. At low densities, nodes are typically too far away to forge competent links. When a new node is added, the $\sigma_r = 0$ model has no preference in choosing m nodes within a quality connection range and thus new nodes fail to connect reliably and the average degree collapses. For $\sigma_r = 1$, there exists some preference for choosing nearby nodes which is capable of slowing the $\langle k \rangle$ decay. Forcing a stronger link length dependence at $\sigma_r = 2$ slows this rate even further and raises the degeneracy with respect to node number.

Clearly, setting the parameter $\sigma_r = 2$ breaks the “scale-freedom” of this model, since the degree behaviour is no longer independent from N . However, this dependence is an inevitable consequence of utilizing quantum communications. As with the classical internet, it is likely that scale-free properties will emerge in quantum networks. However, network engineers will not turn a blind eye to link length and channel quality, making the $\sigma_r = 0$ model unrealistic, rendering larger values of σ_r meaningful and interesting. Throughout this work we focus on $\sigma_{\text{deg}} = 1$, $\sigma_r > 0$ models as better reflections of future quantum networks which we refer to as the network class $\mathcal{S}_{\sigma_r}$.

IV. PRACTICAL ROUTING AND PERFORMANCE-DEFINED CRITICALITY

A. Benchmarking Performance

Consider a class of quantum networks $\mathcal{N} = \{\mathcal{N}_i\}_i$ and a specific routing protocol $\mathcal{P}$. For any network $\mathcal{N} = (P, E) \in \mathcal{N}$ drawn from this class, one can define the average end-to-end rate as that which is averaged over all possible end-user pairs in the network. We denote this set of end-users by

$$\mathcal{I} := \{\mathbf{i}_j\}_j = \{\{\mathbf{a}_j, \mathbf{b}_k\}\}_{j \neq k}. \quad (11)$$

Then the average end-to-end rate takes the form,

$$\langle K \rangle_{\mathcal{N}|\mathcal{P}} := \frac{1}{|\mathcal{I}|} \sum_{\mathbf{i} \in \mathcal{I}} K(\mathbf{i}, \mathcal{N}|\mathcal{P}) \quad (12)$$

where $K(\mathbf{i}, \mathcal{N}|\mathcal{P})$ is the rate achieved between the end-user pair $\mathbf{i} = \{\mathbf{a}, \mathbf{b}\}$ given that they engage in the routing protocol $\mathcal{P}$ (given in Eq. (8)). In an N node network there exist C_N^2 potential end-user pairs (where C_n^k denotes the binomial coefficient), which may be far too large to consider explicitly. In practice, we will sample L end-user pairs from $\mathcal{N}$ to approximate the average end-to-end rate, $\langle K \rangle_{\mathcal{N}|\mathcal{P}} \approx \frac{1}{L} \sum_{j=1}^L K(\mathbf{i}_j, \mathcal{N}|\mathcal{P})$ which can be computed along with statistical error considerations.

This quantity benchmarks the ability of quantum communication on a specific network. To assess the efficacy of a routing strategy more generally, we study the *ensemble average end-to-end rate* sampled across the entire network class. Given that any network $\mathcal{N} \in \mathcal{N}$ is drawn

with equal probability, then the ensemble average rate is equal to

$$\langle K \rangle_{\mathcal{N}|\mathcal{P}} := \frac{1}{|\mathcal{N}|} \sum_{\mathcal{N} \in \mathcal{N}} \langle K \rangle_{\mathcal{N}|\mathcal{P}}. \quad (13)$$

For the network classes with which we are interested, $\mathcal{N}$ can be incredibly large (or infinite). Hence, we are not be able to exactly compute Eq. (13) but instead compute an accurate approximation over a sample space of M' networks $\{\mathcal{N}_i\}_{i=1}^{L'}$ such that $\langle K \rangle_{\mathcal{N}|\mathcal{P}} \approx \frac{1}{L'} \sum_{i=1}^{L'} \langle K \rangle_{\mathcal{N}_i|\mathcal{P}}$. This approximation can be made sufficiently accurate by taking enough end-user samples L and network class samples L' .

These concepts can be immediately translated in the context of end-to-end capacities rather than rates. Indeed, one can readily define a specific end-to-end capacity $\mathcal{C}(\mathbf{i}, \mathcal{N}|\mathcal{P})$, an average $\langle \mathcal{C} \rangle_{\mathcal{N}|\mathcal{P}}$ and ensemble average $\langle \mathcal{C} \rangle_{\mathcal{N}|\mathcal{P}}$. The only difference is the description of the point-to-point rates throughout the network; when network links are considered to operate at their capacity, then we are studying the end-to-end capacities. Otherwise, they are sub-optimal end-to-end rates.

B. Benchmarking Efficiency

Routing protocols should not only be benchmarked with respect to performance, but also with respect to efficiency, i.e. what proportion of network resources is required to guarantee a particular level of performance? Therefore, it is useful to define a measure which we call the *routing consumption*, which quantifies the proportion of network edges used via a given routing protocol. Given a network $\mathcal{N} = (P, E)$, an end-user pair $\mathbf{i} = \{\mathbf{a}, \mathbf{b}\}$ and a routing protocol $\mathcal{P}$, the routing consumption $\tilde{E}$ is the fraction of network edges required to perform communication,

$$\tilde{E}(\mathbf{i}, \mathcal{N}|\mathcal{P}) := \frac{|E_{\mathcal{P}}(\mathbf{i}, \mathcal{N})|}{|E|}, \quad (14)$$

where $|E_{\mathcal{P}}(\mathbf{i}, \mathcal{N})|$ is the number of edges in the routing edge set connecting the end-user pair. We can then define an average end-to-end routing consumption by averaging over end-user pairs $\langle \tilde{E} \rangle_{\mathcal{N}|\mathcal{P}} := \frac{1}{|\mathcal{I}|} \sum_{\mathbf{i} \in \mathcal{I}} \tilde{E}(\mathbf{i}, \mathcal{N}|\mathcal{P})$, which of course extends to an ensemble average over a class of networks $\langle \tilde{E} \rangle_{\mathcal{N}|\mathcal{P}} := \frac{1}{|\mathcal{N}|} \sum_{\mathcal{N} \in \mathcal{N}} \langle \tilde{E} \rangle_{\mathcal{N}|\mathcal{P}}$, following the notation convention used for end-to-end rates. As before, these quantities are estimated by sampling from the set of end-user pairs and the network class.

The routing consumption is a useful measure of how resource efficient a network protocol is at achieving its rates. A protocol which can attain high-rates with a low routing consumption is clearly a desirable strategy. Hence, a tradeoff exists between end-to-end rates and routing consumption. While we know flooding to be rate optimal, it will always satisfy $\langle \tilde{E} \rangle_{\mathcal{N}|\mathcal{P}_{\text{fl}}} = 1$ because it utilises all network edges; hence it is most likely

sub-optimal from the perspective of routing consumption. For single-path protocols we can typically expect $\langle \tilde{E} \rangle_{\mathcal{N}|\mathcal{P}_{\text{sp}}} \ll 1$, but its rates may be poor. It is most interesting to inspect the behaviour of $\langle \tilde{E} \rangle_{\mathcal{N}|\mathcal{P}_{\text{mp}}}$ as understanding the relationship between end-to-end rate, resource consumption and practical multi-path routing is least understood.

C. Practical Routing

In practice, one must be able to deploy a routing algorithm which informs network nodes how to interact and establish communication between the end-users. This algorithm characterises the network protocol, for which there are some intuitive, practical objectives. The algorithm should:

1. Identify an end-to-end path (or paths) which maximize the rate between users. Alternatively, the protocol should surpass a target rate requirement.
2. Minimize the network resources needed to achieve said end-to-end rate, i.e. minimize the necessary network nodal density as well as the number of links/nodes which participate in communication.
3. Be efficient enough to run and facilitate routing without impeding the end-to-end rate.

It is extremely important that each of these points are considered for a routing strategy to be deemed practical. A easily executed protocol is useless if it identifies poor routes, while a protocol which establishes high rate routes very slowly is equally undesirable.

Single-path routing is the principal mechanism for classical communications, fundamentally achieved by Dijkstra's algorithm (DA) [49]. Generally, this is a greedy algorithm for single-path route optimisation according to a path defined *cost function*, which measures a property of end-to-end routes that one wishes to optimise, e.g. minimisation of path length or maximisation of bottleneck rate. The challenges which extends from the point-to-point rate limitations of quantum communications means that rate optimisation is the salient task. This is also known solving the widest path problem. For a network $\mathcal{N} = (P, E)$, DA locates the widest path efficiently in run-time $\mathcal{O}(|E| + |P| \log_2 |P|)$. Given a single-path protocol $\mathcal{P}_{\text{sp}}$ and an end-user pair, the optimal end-to-end rate is given by

$$K(i, \mathcal{N}|\mathcal{P}_{\text{sp}}) := \max_{\omega} \min_{(\mathbf{x}, \mathbf{y}) \in \omega} K_{\mathbf{x}\mathbf{y}}, \quad (15)$$

where the maximisation is performed over all possible end-to-end routes [50].

To mitigate single-path performance limitations, multi-path routing emerges as good solution. Flooding represents a useful paradigm for benchmarking optimal network performance, but it is not practical in many

real world settings. Utilizing an entire network to facilitate one end-user pair renders the network useless for any other set of users to communicate simultaneously. In reality, we want to enable the concurrent use of a network for many users.

A basic approach to building high-rate multi-paths is to perform an iterative form of DA. This is a modified algorithm which locates multiple end-to-end routes which are either edge-disjoint (no two paths share the same edge) or node-disjoint (no two paths visit the same node). Both of these versions involve executing a Dijkstra search, followed by a network modification where edges included in previous paths (or in the neighbourhoods of nodes in previous paths) are not included in the next search. Locating M end-to-end paths requires M executions of DA, leading to an increased time complexity $\mathcal{O}(M(|E| + |P| \log_2 |P|))$. This a costly scaling factor which does not lend well to deployability in large-scale networks.

D. Efficient Multi-Path Routing

Fortunately, there are ongoing developments in the study of fast algorithms for multi-path routing [51–55]. In particular, Lopez-Parajes *et al.* recent devised an efficient, centralised, *one-shot* approach to multi-path routing through their Multiple Disjoint Path Algorithm (MDPAlg) [39]. They recognised that a single execution of DA observes much more information than it actually utilises. Given a source node $\mathbf{a}$, the standard DA focuses on building a minimum cost tree from which only the optimal end-to-end paths can be located from any other node $\mathbf{x} \in P \setminus \{\mathbf{a}\}$ in the network. It does this by storing a minimum cost set (stores the minimum cost associated with traversing from $\mathbf{a}$ to any node $\mathbf{x}$) and a parent node set (stores the parent node from which the minimum cost was obtained to $\mathbf{x}$). With these quantities, DA can then reconstruct a minimum cost path by backtracking from a target node $\mathbf{b}$ along the parent-node tree until $\mathbf{a}$ is reached.

Alternatively, the MDPAlg uses a *cost matrix* to collect additional information on the *aggregated* cost from $\mathbf{a}$ to any other node; not just the minimum cost. Whereas DA would discard information concerning sub-optimal paths, the MDPAlg stores such information in the cost matrix so that they may be used to construct additional end-to-end paths later. Furthermore, it does this through only a single search of the network rather than the many searches that may be required by an iterative Dijkstra approach. The cost matrix can then be used to reconstruct the optimal path, and many other paths between the source $\mathbf{a}$ and $\mathbf{b} \in P \setminus \{\mathbf{a}\}$.

The original algorithm identifies multiple *shortest paths* but can be modified to optimise alternative cost functions. In this work, we modify the MDPAlg in such a way that approximates rate optimisation by minimizing the *inverse accumulated rate* over the course of an end-

to-end route. We may call this variant *IAR-MDPAlg*. We do not minimize the total route length, but rather the sum of its inverted rates. More precisely, we utilise the following optimisation ansatz to locate high-quality, low-resource intensive routes,

$$\omega^* := \arg \min_{\omega \in \Omega} \sum_{(x,y) \in \omega} (K_{xy}^{-\eta} + \epsilon). \quad (16)$$

Here, ω^* is an optimised route, Ω is the set of all possible paths between a given end-user pair. Meanwhile $\eta, \epsilon \in \mathbb{R}_0^+$ are real hyperparameters. This aligns with minimizing the sum of the inverse point-to-point rates along a path ω , adding a penalty for channel usage ϵ associated with each link. It follows that η controls the strength of penalizing the use of links with weak rates, while ϵ controls the penalty strength of resource consumption. Minimizing this cost function simultaneously locates a path which maximizes the sum of the point-to-point rates along the path while minimizing its length; *indirectly* identifying a high-rate and edge-efficient route.

In this way, the algorithm will (typically) identify routes with large bottleneck rates. While this approach is approximate, we show in this work that it is effective (for more details we refer the reader to Appendices C and D). Furthermore, we focus on the edge-disjoint format under the assumption that repeaters may have multiple quantum registers that operate concurrently to route quantum systems along different channels. The edge-disjoint version can locate more end-to-end paths than its node-disjoint counterpart, thus offering greater utility.

There are two potential versions of multi-path protocol that we may consider. To enhance rates, one may insist upon the use of $M > 1$ end-to-end paths, where M is fixed. This is an intuitive approach which we call fixed route number protocols, denoted by $\mathcal{P}_{\text{mdp}}^M$. However, a pair of end-users may be more interested in preserving resources granted that they possess a particular rate guarantee. It is easy to devise a protocol in which there is a rate requirement; via the IAR-MDPAlg, we use the least number of end-to-end routes required to achieve R^* bits per protocol use. Denoted by $\mathcal{P}_{\text{mdp}}^{R^*}$, this is easily implemented strategy and reflects the quality of service principle in classical networks.

It is important to note that the IAR-MDPAlg is not optimal and may not always match the performance of iterative Dijkstra (or, of course, flooding). However, we will see that any minor cost in performance is vastly outweighed by gains in efficiency, as the IAR-MDPAlg can construct end-to-end multi-paths orders of magnitude faster than iterative Dijkstra approaches (see Ref. [39] for cost analyses). This makes it a much more practical method and allows us to perform numerical assessments that would not be possible otherwise. For a more detailed explanation of the algorithms used in this study, we refer the reader to Appendix D.

E. Criticality of Quantum Networks

It is essential to pursue a rigorous and quantitative assessment of practicality with respect to quantum network routing, and allow us to understand the core network features which guarantee both performance and efficiency. To this end, the concept of *network criticality* is vital. A critical transition is an abrupt regime shift in the behaviour of some complex system, such that the system transits from sub-to-super-critical with respect to a particular property. In the context of complex quantum networks, we have already discussed how critical transitions may occur with regards to robustness (connectivity) but it can be extended to notions of performance [34] or routing efficiency. A major contribution of our work is to introduce routing consumption-criticality, identifying network regimes within which end-to-end routing is efficient via practical routing.

We begin with performance-defined criticality; critical properties within a quantum network model which a guarantee a transition from unreliable rates to consistent super-critical rates. A useful manifestation of performance-based criticality that can arise is end-to-end distance independent rates, i.e. the spatial separation of users is independent of communication quality. This is especially important to quantum networks in order to overcome point-to-point limitations. Furthermore, these measures are more appropriate than purely connectivity defined quantities, since they make meaningful assessments of the efficacy of a network to perform quantum communications. This quantity was first introduced in [34] in the context of network flooding protocols $\mathcal{P}_{\text{fl}}$ and was used to illuminate critical insight for quantum network composed of bosonic pure-loss channels. In this work, we explore such critical densities relative to a broader range of routing protocols and link-layer descriptions.

Definition 1 (Performance-defined Critical Density): *For a class of quantum networks $\mathbf{N}$, and a network routing protocol $\mathcal{P}$, we define $\rho_{\mathbf{N}|\mathcal{P}}^*$ as the minimum nodal density required to guarantee an ensemble average rate of $\langle R \rangle_{\mathbf{N}|\mathcal{P}} \geq 1$ bit/protocol use.*

The critical nodal density $\rho_{\mathbf{N}|\mathcal{P}}^*$ is an extremely important measure since quantum repeaters are costly and they should be minimized as a resource in a future quantum internet. Investigations of this quantity with respect to flooding $\rho_{\mathbf{N}|\mathcal{P}_{\text{fl}}}^*$ and capacity achieving bosonic lossy networks have been carried out, which can be considered a lower-bound for all other critical densities. Indeed, minimizing this quantity over all protocols is equivalent to the flooding defined critical density,

$$\rho_{\mathbf{N}}^* := \min_{\mathcal{P}} \rho_{\mathbf{N}|\mathcal{P}}^* = \rho_{\mathbf{N}|\mathcal{P}_{\text{fl}}}^* \leq \rho_{\mathbf{N}|\mathcal{P}}^*. \quad (17)$$

Nonetheless, analogous studies have not been carried out for other routing protocols or networks.

The performance-based critical density is a valuable characteristic of a quantum network model, however it lacks insight to the resources required to achieve critical rates. Networks may be super-critical with respect to performance but may demand impractical resources to do so. To address this we can analyse efficiency regimes with respect to routing, defining a consumption-based critical density.

Definition 2 (Consumption-defined Critical Density): *Consider a class of quantum networks $\mathbf{N}$, and let $\mathbf{N}(\rho) \subset \mathbf{N}$ be a subset of this class with nodal density ρ . For a network routing protocol $\mathcal{P}$, we define $\tilde{\rho}_{\mathbf{N}|\mathcal{P}}^*$ as the density at which the routing consumption is maximized and after which it undergoes consistent decay. More precisely, $\tilde{\rho}_{\mathbf{N}|\mathcal{P}}^* := \arg \max_{\rho} \langle \tilde{E} \rangle_{\mathbf{N}(\rho)|\mathcal{P}}$, such that $\langle \tilde{E} \rangle_{\mathbf{N}(\rho)|\mathcal{P}} \leq \langle \tilde{E} \rangle_{\mathbf{N}(\tilde{\rho}_{\mathbf{N}}^*)|\mathcal{P}}$, $\forall \rho \geq \tilde{\rho}_{\mathbf{N}}^*$.*

Hence, $\tilde{\rho}_{\mathbf{N}|\mathcal{P}}^*$ represents a secondary critical measure which separates network classes into different efficiency regimes. Executing the same protocol, networks which possess a nodal density $\rho > \tilde{\rho}_{\mathbf{N}|\mathcal{P}}^*$ will consume a smaller fraction of network resources during routing. Furthermore, this efficiency increase as the density continues to grow. We may connect the concepts of performance and routing consumption within the following definition.

Definition 3 (δ -Critical routing consumption): *Consider a class of quantum networks $\mathbf{N}$. The critical routing consumption $\langle \tilde{E} \rangle_{\mathbf{N}}^*$ is the minimum ensemble average fraction of network edges that must engage in end-to-end routing granted the ensemble average rate is $\langle R \rangle_{\mathbf{N}|\mathcal{P}} \geq \delta$ bits per network use, given that $\delta \leq \langle R \rangle_{\mathbf{N}|\mathcal{P}_{\text{R}}}$. More precisely, the critical routing consumption is*

$$\langle \tilde{E} \rangle_{\mathbf{N},\delta}^* := \min_{\mathcal{P}: \langle R \rangle_{\mathbf{N}|\mathcal{P}} \geq \delta} \langle \tilde{E} \rangle_{\mathbf{N}|\mathcal{P}}. \quad (18)$$

The definition in Eq. (18) can be made intuitive in the following way: There exists a critical routing consumption achieved by an end-to-end protocol which can promise δ bits per network use on average, such that δ is an attainable network rate. Further to this performance guarantee, the protocol minimizes the ensemble average fraction of network edges used, i.e. on average, it is the most efficient routing mechanism which promises the target rate. It is by no means clear how to determine the optimal protocol. However, the tools developed in this paper can effectively bound the critical routing consumption via practical routing schemes. In fact, we may write

$$\langle \tilde{E} \rangle_{\mathbf{N}|\mathcal{P}_{\text{sp}}} \leq \langle \tilde{E} \rangle_{\mathbf{N},\delta}^* \leq \langle \tilde{E} \rangle_{\mathbf{N}|\mathcal{P}_{\text{mp}}}. \quad (19)$$

The optimal single-path protocol can always be a lower-bound, which is saturated iff an ensemble average δ bits per network use is achievable via single-path routing. An appropriate multi-path strategy generates the upper-bound. Indeed, there will always exist a multi-path strategy capable of this since one can employ a flooding protocol in the worst case (least efficient) scenario.

V. NUMERICAL RESULTS

A. Benchmarking Waxman Fiber Networks

In Fig. 3 we present relationships between nodal density, routing consumption and end-to-end capacities on bosonic thermal-loss Waxman networks, with an upper-bounding capacity distribution $\mathcal{K}_u$ (lower-bounding capacity distribution $\mathcal{K}_l$). Each network edge represents an optical-fiber link of loss rate 0.2 dB/km and environmental thermal noise $\bar{n} = 1/500$ connecting ideal transmitter/detectors. The Waxman parameters are set as $r_0, \beta = 100, 1$ ($r_0, \beta = 63, 1$), where the decay parameter corresponds to the approximate distance at which the single-edge capacity upper-bound collapses to zero ~ 100 km (~ 63 km). Since we are considering end-to-end *capacities*, results concerning these networks offer universal benchmarks for any fiber-based Waxman network.

Figs. 3(a)-(b) depict the ensemble average end-to-end capacities with respect to nodal density for a number of routing strategies. In Ref. [34], the authors identified a performance-based critical nodal density of $\rho_{\mathbf{N}}^* \approx 4.25 \times 10^{-4}$ nodes/km² associated with flooding and bosonic pure-loss networks. The consideration of environmental thermal-noise naturally increases the critical density such that $\rho_{\mathbf{N}}^* \approx 7.35 \times 10^{-4}$ (1.04×10^{-3}) nodes/km² (the density in parentheses refers to the lower-bounding capacity distribution). While significant, this shift remains relatively optimistic since the considered thermal-loss link layer does not consider additional experimental sources of thermal noise in an effort to remain protocol agnostic.

Critical densities with respect to flooding protocols are inherently optimistic, due to the unrealistic resource demand. This optimism is emphasised when one considers the utility of single-path routing. Fig. 3(a) clearly illustrates the intense demands on the network density required for single-path routing to reach criticality, and promise effective end-to-end rates. We find the single-path critical nodal density to be approximately 8.52×10^{-3} (8.86×10^{-3}) nodes/km², which is an order of magnitude larger than that predicted by flooding. On large-scales, a stark increase of this magnitude has significant ramifications for the cost and deployability of quantum networks. Therein lies the necessity for efficient multi-path routing protocols. Fig. 3(b) illustrates the efficacy of IAR-MDPAI based routing, showing that flooding is not necessary to preserve lower critical densities. Each variant of the multi-path protocol is able to recover a critical density close to that offered by flooding, reinforcing the notion that high-rates can be guaranteed with realistic resources.

Further evidence of this is gathered in Figs. 3(c)-(d) which plot the routing consumption of the protocols $\mathcal{P}_{\text{sp}}$ and $\mathcal{P}_{\text{mdp}}^{R^*=1}$ respectively. While single-path protocols are expected to achieve very low routing consumptions, the same expectation is not necessarily held for multi-path routing. However, the protocol $\mathcal{P}_{\text{mdp}}^{R^*=1}$ can

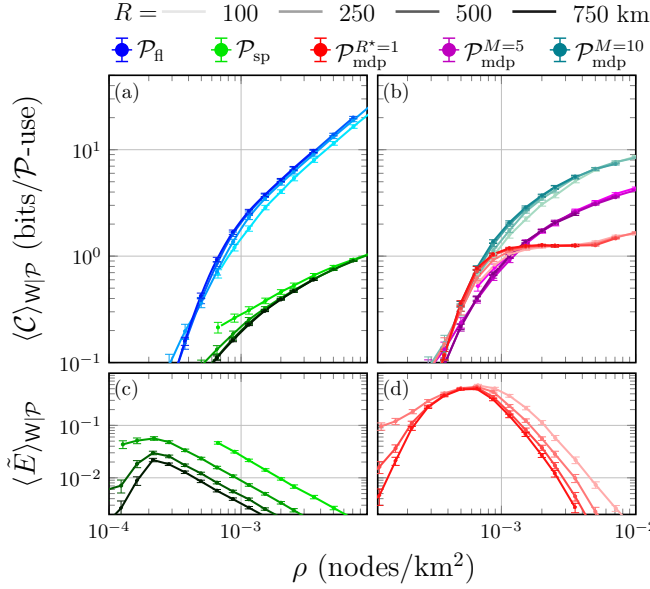


Figure 3. Relationships between performance, routing consumption and nodal density in bosonic thermal-loss Waxman networks. The network link layer in all cases is described by the upper-bound capacity distribution $\mathcal{K}_u$ from Eq. (4) and model parameters $(r_0, \beta) = (100, 1)$. Panels (a)-(b) plot the ensemble average end-to-end capacities achieved by different routing protocols and a number of network radii, R (each protocol and network radius is colour coded with the routing protocol and R legends). Panels (c)-(d) show the ensemble average routing consumption of (c) single path routing and (d) $\mathcal{P}_{mdp}^{R^*=1}$ routing via the IAR-MDPAIlg. Panel (e) depicts achievable end-to-end routes on an example network (according to the parameters listed) using $\mathcal{P}_{sp}$ and $\mathcal{P}_{mdp}^{R^*=1}$ for users separated by $r_i \approx 800$ km. Black edges in the networks identify unused edges, while red edges are those which engage in the routing protocol.

obtain a critical density of approximately 8.15×10^{-4} (1.21×10^{-3}) nodes/km², while maintaining a average maximum routing consumption of approximately 0.44 (0.36). That is, no more than half of the network edges are ever required to participate in end-to-end routing. Even before single-path routing is a plausible option, the routing consumption of $\mathcal{P}_{mdp}^{R^*=1}$ undergoes a rapid decay as the nodal density passes approximately 6.59×10^{-4} (1.15×10^{-3}) nodes/km². For networks with nodal density $\rho \geq 3 \times 10^{-3}$, this protocol can guarantee super-critical rates while consuming less than 5% of the network edges during routing.

In this way, we can use the multi-path protocol $\mathcal{P}_{mdp}^{R^*=1}$ to upper-bound the $\delta = 1$ critical routing consumption as defined in Eq. (19), i.e. for bosonic thermal-loss networks we can write an upper-bound on the minimum network fraction required to guarantee an ensemble average of 1 bit/protocol-use, $\langle \tilde{E} \rangle_{W,1}^* \lesssim 0.44$. This places an efficiency bound on the optimal end-to-end network protocol, which is significantly more informative than flooding.

Minimizing edge usage while guaranteeing high rates is vital to scalability in large quantum networks with many users. Crucially, in the absence of effective single-path routing, efficient multi-path strategies exist that preserve network resources while obtaining reliable, super-critical end-to-end rates. The importance of this result is emphasized when more realistic links are considered, as those described by practical QKD protocols (and beyond).

B. Practical Link Layers and Network Phases

Performing the previous analyses of end-to-end rates, routing consumption, and nodal density for different link layers, we can build a comprehensive picture of the efficacy of quantum communication networks in different scenarios. As such, we can establish relationships between different critical network properties which give rise to key *network phases*; sub-classes of Waxman networks for which the behaviour of end-to-end quantum communication have similar characteristics. Here, phases are defined through statistical analyses of the properties deemed most important to end-to-end communication; connectivity, performance, and routing efficiency. Identifying what these phases are and where they fall within meaningful density ranges can help us to understand the realistic needs of quantum networking.

Fig. 4 summarises the Waxman quantum network phases defined in this study for a number of link layer models: bosonic thermal-loss capacity distributions, and asymptotic CV-QKD rate distributions. Here, we have identified six key network phases, explicitly defined in Fig. 4(b), ranging from Phase I networks with no connectivity guarantees, to Phase VII networks which promise super-critical performance via single-path routing. In between, there exists a spectrum of phases corresponding to critical changes in the connectivity, routing consumption and performance guarantees.

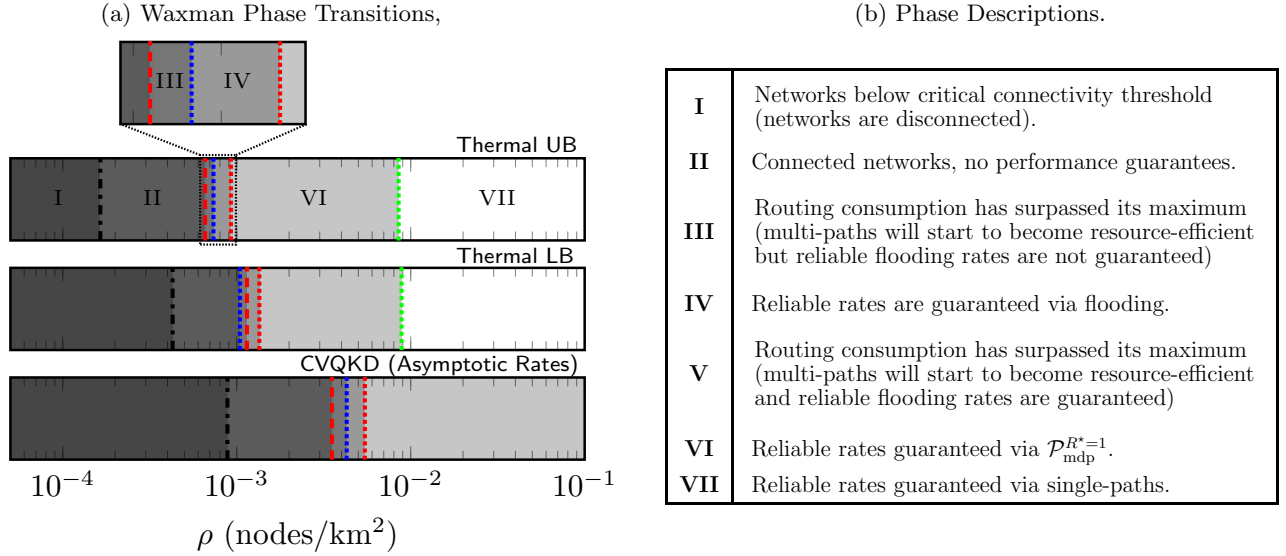


Figure 4. Waxman quantum network phase characterisations with respect to link layer descriptions and nodal density. Panel (a) outlines connectivity, consumption and performance based critical densities with respect to networks composed of different link layers. These transitions give rise to network phases in which we can expect particular properties. These phases are labelled on the density diagram, while Panel (b) summarises and describes their implications for quantum networking. Phases III and V describe similar network properties, but differ as to whether the maximum routing consumption is surpassed before or after the flooding-based performance transition. Note that Phase III appears in networks described by the thermal upper-bound capacity distribution, while Phase V appears in the lower-bound capacity distribution.

These phases have been identified by computing and comparing various critical densities relative for each link layer model. Let us define these phases more precisely:

- **I:** Densities fall below the giant-component transition density, $\rho < \rho_G^*$.
- **II:** Densities are greater than the giant-component transition density, but smaller than the flooding-based critical density, $\rho_G^* \leq \rho \leq \rho_{W|\mathcal{P}_{\text{fi}}}^*$.
- **III:** Densities are greater than the consumption-defined critical density, but smaller than the flooding-based critical density, $\tilde{\rho}_{W|\mathcal{P}_{\text{fi}}}^* \leq \rho \leq \rho_{W|\mathcal{P}_{\text{fi}}}^*$. Not all link-layer models will inhabit this phase.
- **IV:** Densities are greater than the flooding-based critical density, $\rho \geq \rho_{W|\mathcal{P}_{\text{fi}}}^*$.
- **V:** Densities are greater than the flooding-based critical density and greater than the consumption-defined critical density $\rho \geq \tilde{\rho}_{W|\mathcal{P}_{\text{fi}}}^* \geq \rho_{W|\mathcal{P}_{\text{fi}}}^*$. Not all link-layer models will inhabit this phase.
- **VI:** Densities are greater than the $\mathcal{P}_{\text{mdp}}^{R^*=1}$ -based critical density, $\rho \geq \rho_{W|\mathcal{P}_{\text{mdp}}^{R^*=1}}^*$.
- **VII:** Densities are greater than the single-path-based critical density, $\rho \geq \rho_{W|\mathcal{P}_{\text{sp}}}^*$.

Fig. 4(b) summarises these phases and interprets them from a practical perspective.

The most desirable network phase is naturally Phase VII, in which single-path routing is sufficient to perform reliable quantum communication. Unfortunately, our primary takeaway is that the nodal densities necessary to inhabit Phase VII are very large, and become impractical as point-to-point link layer descriptions become more realistic. Even when considering fully trusted QKD networks described using asymptotic secret key-rates, the nodal densities required to reach Phase VII are more than 10^{-1} nodes/km². For example, the number of nodes needed to deploy a fully trusted QKD network that spans the surface area of Europe ($\sim 10^7$ km²) would be on the order of millions.

Fortunately, efficient multi-path routing strategies can resurrect the utility of lower density networks, which achieve reliable rates. Networks which occupy Phase VI are super-critical provided that they employ the protocol $\mathcal{P}_{\text{mdp}}^{R^*=1}$, or better. As shown in Fig. 4 for each considered link layer, Phase VI occupies a more practical nodal density region (around an order of magnitude improvement) for which practical multi-path strategies can promise strong rates using efficient protocols; not just flooding.

An important observation that we wish to reiterate is the existence of Phases II and VI. Phase II portrays the resource gap between connectivity guarantees and the most optimistic performance guarantee (via flooding). The existence of this gap makes it clear that the design and assessment of quantum networks *cannot* solely focus on connectivity analyses, as such promises are not enough to guarantee useful quantum communication. Analo-

gously, Phase VI identifies a gap between the resources required by flooding and practical multi-path protocols for critical rates. Closing the gap posed between Phase VI and VII while minimizing routing consumption is the goal of any multi-path strategy.

These insights emphasise the need for explicit analyses of end-to-end performance and reiterate the point that routing in quantum networks cannot naïvely follow its classical counterpart. Multi-path routing techniques do not just represent a means of boosting rates but establish a pertinent method of reducing the resource demands of practical quantum network development. We reiterate that these phases are specifically derived for the Waxman network model, which is defined by its relationship between link-length and network connectivity. Nonetheless, the rate-distance tradeoff is a quintessential element of quantum communications; hence, these results provide an informative window into future quantum network properties.

C. Scale-Free Properties and Quantum Networks

Scale-free architectures capture important features of real world networks beyond that of the Waxman model. As discussed in Section III C, scale-free networks obey a power-law cumulative degree distribution. This realises a connectivity structure in which there exist a number of highly connected hubs to which many nodes of low degree are connected, giving rise to a lower average degree than Waxman models. Considering the network class $\mathcal{S}_{\sigma_r}$, we can investigate the ramifications these connectivity features have on end-to-end routing and performance

Fig. 5 displays analyses of bosonic thermal-loss scale-free networks with respect to nodal density and a number of routing protocols: Flooding, single-path, and $\mathcal{P}_{\text{mdp}}^{R^*=1}$ routing using the edge-disjoint IAR-MDPAlg. We consider architectures generated with the parameter $\sigma_r \in \{1, 2\}$ so to vary the awareness of link-length during network creation and consequently its “strictness” with respect to scale-free behaviour. In Figs. 5 (a)-(c) we study the ensemble average end-to-end capacities for each network model and routing strategy. It is clear for $\sigma_r = 1$ (in which scale-free behaviour is followed closely) that performance is restricted by the low average degree. At small network scales (e.g. $R < 150$ km) and increasing density, the flooding capacity is able to reach reliable rates, but only at very high densities. Furthermore, Ref. [34] showed that the optimal flooding capacity of such networks exponentially decays with respect to R (regardless of the number of nodes).

Practical routing strategies cannot do better, as displayed for both single-path routing and $\mathcal{P}_{\text{mdp}}^{R^*=1}$ routing. Ultimately, the connectivity structure of scale-free networks $\mathcal{N} \in \mathcal{S}_1$ display a poor aptitude for multi-path routing. The utilisation of multiple routes is only helpful if many additional routes can be found, and the existence of high-degree hubs limits this possibility. User

nodes will typically only connect to a single hub, limiting the ability to reinforce the end-to-end path set. This limitation can be seen in the ensemble average routing consumption of the $\mathcal{P}_{\text{mdp}}^{R^*=1}$ protocol, which remains relatively constant with respect to nodal density. This means that there is little correlation between increasing density and the number of effective end-to-end routes between user nodes.

Breaking from strict scale-freedom, we see that effective performance *can* re-emerge within the network class $\mathcal{S}_2$. With an increased connection probability dependence on link-length, the ensemble average capacity and routing consumption follow similar trends to the Waxman model: The exponential decay of the flooding capacity with respect to network scale is subdued, and end-to-end capacities of $\langle \mathcal{C} \rangle_{\mathcal{S}_2|\mathcal{P}} \geq 1$ bit/protocol-use can be guaranteed for each routing protocol (within reasonable nodal density ranges). Furthermore, the multi-path routing consumption reliably decays with increasing density, implying that an applicability for multi-path routing can be reestablished.

VI. DISCUSSION

In this work, we have investigated the capabilities of quantum communication networks under practical routing strategies, in an effort to gain insight to realistic requirements of future quantum networks. With the goal of providing a comprehensive study of the performance and feasibility of quantum optical fiber-networks, we have combined theory of quantum communication with that from random network models and network routing algorithms. Our work departs from previous developments in this domain by considering realistic link layer descriptions (capacities and secret-key rates over bosonic thermal-loss channels) alongside practical routing protocols (efficient multi-path and single-path routing methods). These assessments focus on the crucial Waxman and scale-free classes of random quantum networks.

Through our statistical analyses, we reveal vital insights into the practical requirements of critical behaviour in quantum networks, where criticality may be defined with respect to connectivity, performance and routing efficiency. For quantum Waxman networks, a network phase structure is introduced with respect to nodal density; identifying ranges of nodal densities within which networks can be classified as sub- or super-critical with respect to these properties. We show that for increasingly realistic link-layer descriptions, the existence of a performance super-critical phase with respect to single-path routing becomes more and more difficult within practical density ranges. Nonetheless, multi-path routing strategies can re-establish practical density ranges and super-critical rates, reiterating the need for the development of multi-path routing methods in quantum networks.

It is important to note that the network phases iden-

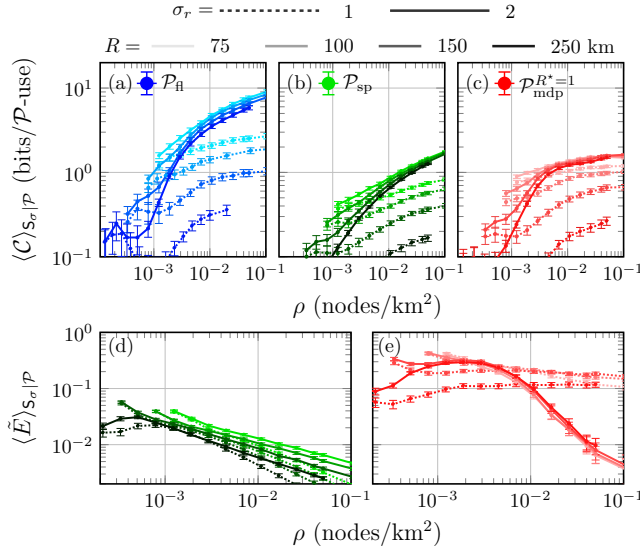


Figure 5. Relationships between performance, routing consumption and nodal density in bosonic thermal-loss scale-free networks. Throughout all plots, the network link layer is described by the upper-bound capacity distribution $\mathcal{K}_u$ from Eq. (4) and we use the model parameters $(n_0, m, \sigma_{\text{deg}}) = (10, 5, 1)$. We also consider unique values for the scale-free model parameter $\sigma_r \in \{1, 2\}$ which are distinguished in the legend. Panels (a)-(c) depict the ensemble average capacity with respect to nodal density for $\mathcal{P}_R$, $\mathcal{P}_{sp}$ and $\mathcal{P}_{mdp}^{R^*=1}$ routing respectively. Panels (d)-(e) plot the ensemble average routing consumption with respect to nodal density for $\mathcal{P}_{sp}$ and $\mathcal{P}_{mdp}^{R^*=1}$ routing respectively. Each analysis is performed for a number of network radii R listed in the legend. Panel (f) illustrates achievable end-to-end routes on an example network (under the parameters shown) using $\mathcal{P}_{sp}$ and $\mathcal{P}_{mdp}^{R^*=1}$ for users separated by $r_i \approx 200$ km. Black edges in the networks identify unused edges, while red edges are those which engage in the routing protocol.

tified are classifications which emerge naturally from the heuristics considered in this work but are by no means definitive. With future, increasingly sophisticated investigations we expect richer and more detailed phases to be unveiled which provide insight and guidance for quantum network design.

Our findings reiterate the challenging relationship between scale-free networks and quantum communication. While scale-free architectures such as Yook’s model [47] offer useful insight into the structure of the classical internet, they do not find an analogy with reliable quantum fiber networks. Adherence to scale-freedom in large-scale quantum networks results in a poor communication quality regardless of routing strategy. By altering Yook’s model to generate networks with an increased sensitivity to link length, scale-freedom begins to break while the ability to perform quantum communication strengthens. These results stress the resounding differences between the classical and quantum internet: Quantum networks need to be *consistently* well connected to overcome point-to-point limitations, and must be strictly engineered with nodal density and link length in mind. This is a crucial take-home message that our work has established.

While our results focus on the efficacy of fiber networks, it is well motivated that the quantum internet will exploit satellite links and inter-satellite networks to reinforce long-range communication. This is an area of serious interest and a future investigative path for ex-

tending this research, investigating the interplay of realistic, ground-based networks interacting with satellite-based infrastructure. Furthermore, the study of ground-based free-space networks in both long-range (inter-metropolitan) and mobile (metropolitan) settings is of immediate interest.

Understanding how many end-user pairs may simultaneously perform quantum communication across realistic networks is of paramount importance. Indeed, the trade-off between network architecture, end-to-end rates, and the number of end-users is still poorly understood. While the study of end-to-end routing consumption is a useful step in this direction, greater progress must be made in order to develop practical routing strategies for the quantum internet. In particular, future work should extend these results to the case of entanglement distribution via imperfect quantum repeaters, where decoherence, operational errors, and finite memory lifetimes will further impact network performance and routing efficiency.

ACKNOWLEDGMENTS

This work was supported by the EPSRC (Grant No. EP/R513386/1) and UKRI through the Integrated Quantum Networks (IQN) Research Hub (Grant No. EP/Z533208/1). C.H thanks Alasdair I. Fletcher for helpful comments and discussions.

-
- [1] H. J. Kimble, “The quantum internet,” *Nature* **453**, 1023 (2008).
 - [2] S. Pirandola and S. L. Braunstein, “Physics: Unite to build a quantum internet,” *Nature* **532**, 169 (2016).
 - [3] M. Razavi, *An Introduction to Quantum Communications Networks* (Morgan & Claypool, 2018).
 - [4] S. Pirandola, U. L. Andersen, L. Banchi, M. Berta, D. Bunandar, R. Colbeck, D. Englund, T. Gehring, C. Lupo, C. Ottaviani, et al., “Advances in quantum cryptography,” *Advances in Optics and Photonics* **12**, 1012 (2020).
 - [5] Y.-A. Chen, Q. Zhang, T.-Y. Chen, W.-Q. Cai, S.-K. Liao, J. Zhang, K. Chen, J. Yin, J.-G. Ren, Z. Chen, S.-L. Han, et al., “An integrated space-to-ground quantum communication network over 4,600 kilometres,” *Nature* **589**, 214 (2021).
 - [6] J. S. Sidhu, S. K. Joshi, M. Gundogan, T. Brougham, D. Lowndes, L. Mazzarella, M. Krutzik, S. Mohapatra, D. Dequal, G. Vallone, et al., “Advances in space quantum communications,” *IET Quantum Communication*, 1 (2021).
 - [7] J. I. Cirac, A. K. Ekert, S. F. Huelga, and C. Macchiavello, “Distributed quantum computation over noisy channels,” *Phys. Rev. A* **59**, 4249 (1999).
 - [8] D. Gottesman, T. Jennewein, and S. Croke, “Longer-baseline telescopes using quantum repeaters,” *Phys. Rev. Lett.* **109**, 070503 (2012).
 - [9] P. Kómár, E. M. Kessler, M. Bishof, L. Jiang, A. S. Sørensen, J. Ye, and M. D. Lukin, “A quantum network of clocks,” *Nature Physics* **10**, 582 (2014).
 - [10] M. A. Nielsen and I. L. Chuang, *Quantum Computation and Quantum Information: 10th Anniversary Edition*, 10th ed. (Cambridge University Press, 2011).
 - [11] J. Watrous, *The Theory of Quantum Information* (Cambridge University Press, 2018).
 - [12] A. S. Holevo, *Quantum Systems, Channels, Information* (De Gruyter, 2019).
 - [13] P. Slepian, *Mathematical Foundations of Network Analysis* (Springer-Verlag, New York, 1968).
 - [14] R. Pastor-Satorras and A. Vespignani, *Evolution and Structure of the Internet: A Statistical Physics Approach* (Cambridge University Press, 2004).
 - [15] T. M. Cover and J. A. Thomas, *Elements of Information Theory* (Wiley, New Jersey, 2006).
 - [16] A. S. Tanenbaum and D. J. Wetherall, *Computer Networks*, 5th ed. (Pearson, 2010).
 - [17] A. El Gamal and Y.-H. Kim, *Network Information Theory* (Cambridge University Press, 2011).
 - [18] S. Pirandola, R. García-Patrón, S. L. Braunstein, and S. Lloyd, “Direct and reverse secret-key capacities of a quantum channel,” *Phys. Rev. Lett.* **102**, 050503 (2009).
 - [19] S. Pirandola, R. Laurenza, C. Ottaviani, and L. Banchi, “Fundamental limits of repeaterless quantum communications,” *Nature Communications* **8**, 15043 (2017).
 - [20] S. Pirandola, “Capacities of repeater-assisted quantum communications,” arXiv:1601.00966 (2016).
 - [21] S. Pirandola, “End-to-end capacities of a quantum communication network,” *Communications Physics* **2**, 51 (2019).
 - [22] M. Pant, H. Krovi, D. Towsley, L. Tassioulas, L. Jiang, P. Basu, D. Englund, and S. Guha, “Routing entanglement in the quantum internet,” *npj Quantum Information* **5**, 25 (2019).
 - [23] X. Meng, J. Gao, and S. Havlin, “Concurrence percolation in quantum networks,” *Phys. Rev. Lett.* **126**, 170501 (2021).
 - [24] X. Meng, et al., “Path percolation in quantum communication networks,” *Phys. Rev. Lett.* **134**, 030803 (2025).
 - [25] M. Caleffi, “Optimal routing for quantum networks,” *IEEE Access* **5**, 22299–22312 (2017).
 - [26] S. Shi and C. Qian, “Concurrent entanglement routing for quantum networks: Model and designs,” *Proceedings of the Annual Conference of the ACM Special Interest Group on Data Communication on the Applications, Technologies, Architectures, and Protocols for Computer Communication*, 62–75 (2020).
 - [27] A. Fischer and D. Towsley, “Distributing graph states across quantum networks,” 2021 IEEE International Conference on Quantum Computing and Engineering (QCE), 168–178 (2021).
 - [28] A. Kolar, et al., “Adaptive, continuous entanglement generation for quantum networks,” IEEE INFOCOM 2022 – IEEE Conference on Computer Communications Workshops (INFOCOM WKSHPS), 307–312 (2022).
 - [29] Y. Huang, et al., “Peer-to-peer distribution of graph states across spacetime quantum networks of arbitrary topology,” *Proceedings of the ACM on Measurement and Analysis of Computing Systems* **9**, 1–36 (2025).
 - [30] C. Harney and S. Pirandola, “Analytical methods for high-rate global quantum networks,” *PRX Quantum* **3**, 010349 (2022).
 - [31] C. Harney, A. I. Fletcher, and S. Pirandola, “End-to-end capacities of hybrid quantum networks,” *Phys. Rev. Applied* **18**, 014012 (2022).
 - [32] C. Harney and S. Pirandola, “End-to-end capacities of imperfect-repeater quantum networks,” *Quantum Science and Technology* **7**, 045009 (2022).
 - [33] N. R. Solomons, A. I. Fletcher, D. Aktas, N. Venkatachalam, S. Wengerowsky, M. Lončarić, S. P. Neumann, B. Liu, Ž. Samec, M. Stipčević, R. Ursin, et al., “Scalable authentication and optimal flooding in a quantum network,” *PRX Quantum* **3**, 020311 (2022).
 - [34] Q. Zhuang and B. Zhang, “Quantum communication capacity transition of complex quantum networks,” *Phys. Rev. A* **104**, 022608 (2021).
 - [35] B. Zhang and Q. Zhuang, “Quantum internet under random breakdowns and intentional attacks,” *Quantum Science and Technology* **6**, 045007 (2021).
 - [36] R. Albert and A.-L. Barabási, “Statistical mechanics of complex networks,” *Rev. Mod. Phys.* **74**, 47 (2002).
 - [37] B. Bollobás, R. Kozma, and D. Miklós, *Handbook of Large-Scale Random Networks* (Springer, Berlin Heidelberg, 2008).
 - [38] S. Brito, A. Canabarro, R. Chaves, and D. Cavalcanti, “Statistical properties of the quantum internet,” *Phys. Rev. Lett.* **124**, 210501 (2020).
 - [39] D. Lopez-Pajares, E. Rojas, J. A. Carral, I. Martinez-Yelmo, and J. Alvarez-Horcajo, “The disjoint multipath challenge: Multiple disjoint paths guaranteeing scalability,” *IEEE Access* **9**, 74422 (2021).
 - [40] Note that, in continuous-variable (CV) QKD, thermal noise is used to model fiber imperfections that are not related to loss [41]. In fact, it is typical to introduce the

- channel's excess noise as $\varepsilon := (1 - \eta)\eta^{-1}\bar{n}_{\text{env}}$ [42]. In experimental implementations, the typical value of the excess noise ε is in the range 10^{-3} – 10^{-1} . For a transmissivity of $\eta = 1/2$ (corresponding to 15 km in standard fiber), we then have $\bar{n}_{\text{env}}$ in the same range of 10^{-3} – 10^{-1} .
- [41] F. Grosshans, G. Van Assche, J. Wenger, R. Brouri, N. J. Cerf, and P. Grangier, "Quantum key distribution using gaussian-modulated coherent states," *Nature* **421**, 238–241 (2003).
- [42] S. Pirandola, S. L. Braunstein, R. Laurenza, C. Ottaviani, T. P. W. Cope, G. Spedalieri, and L. Bianchi, "Theory of channel simulation and bounds for private communication," *Quantum Science and Technology* **3**, 035009 (2018).
- [43] It is important to make a distinction between *physical* and *logical* information flow. The logical direction of a quantum communication channel refers to the direction in which entanglement, secret keys, or quantum states are exchanged from node $\mathbf{x}$ to node $\mathbf{y}$ (or vice versa). Physical flow refers to the direction of quantum system exchange, i.e., if quantum systems are physically sent from a node $\mathbf{x}$ to a node $\mathbf{y}$. In general, a quantum channel $\mathcal{E}_{\mathbf{x} \rightarrow \mathbf{y}}$ that physically transmits quantum systems from $\mathbf{x}$ to $\mathbf{y}$ may be different from its reverse channel $\mathcal{E}_{\mathbf{y} \rightarrow \mathbf{x}}$. Thanks to quantum teleportation, either direction of logical information flow can be achieved by using any physical channel direction; hence, it is always possible to choose the *best* physically directed channel for logical quantum communication (e.g., that which has the largest capacity).
- [44] This is already a lenient condition, as clock rates in many quantum communication protocols are thus far restricted to the MHz range. Furthermore, the collapse of the point-to-point thermal loss rate functions means that the minimum length of pruned edges is very similar for threshold values as large as $\varepsilon \sim 10^{-6}$, so this choice is quite flexible.
- [45] The number of nodes in the giant component is less than the total number of network nodes N .
- [46] A. D. Broido and A. Clauset, "Scale-free networks are rare," *Nature Communications* **10**, 1017 (2019).
- [47] S.-H. Yook, H. Jeong, and A.-L. Barabási, "Modeling the internet's large-scale topology," *Proceedings of the National Academy of Sciences* **99**, 13382 (2002).
- [48] A.-L. Barabási and R. Albert, "Emergence of scaling in random networks," *Science* **286**, 509 (1999).
- [49] E. W. Dijkstra, "A note on two problems in connexion with graphs," *Numerische Mathematik* **1**, 269 (1959).
- [50] This is equivalent to specifying Eq. (8) with the forwarding probability distribution to $\{q_{\mathbf{x}\mathbf{y}} = 1\}_{(\mathbf{x}, \mathbf{y}) \in \omega^* \cup \{q_{\mathbf{x}\mathbf{y}} = 0\}_{(\mathbf{x}, \mathbf{y}) \notin \omega^*}$.
- [51] J. O. Abe, H. A. Mantar, and A. Yayimli, " k -maximally disjoint path routing algorithms for SDN," 2015 International Conference on Cyber-Enabled Distributed Computing and Knowledge Discovery, 499 (2015).
- [52] J. Tapolcai, G. Rétvári, P. Babarczi, and E. R. Bérczi-Kovács, "Scalable and efficient multipath routing via redundant trees," *IEEE Journal on Selected Areas in Communications* **37**, 982 (2019).
- [53] O. Lemeshko, O. Yeremenko, M. Yevdokymenko, and B. Sleiman, "Enhanced solution of the disjoint paths set calculation for secure QoS routing," in *2019 IEEE International Conference on Advanced Trends in Information Theory (ATIT)*, 210–213 (2019).
- [54] B. Martín, A. Sánchez, C. Beltrán-Royo, and A. Duarte, "Solving the edge-disjoint paths problem using a two-stage method," *International Transactions in Operational Research* **27**, 435 (2020).
- [55] D. Lopez-Pajares, J. Alvarez-Horcajo, E. Rojas, J. A. Carral, and I. Martínez-Yelmo, "One-shot multiple disjoint path discovery protocol (1S-MDP)," *IEEE Communications Letters* **24**, 1660 (2020).
- [56] S. Pirandola, "Composable security for continuous variable quantum key distribution: Trust levels and practical key rates in wired and wireless networks," *Phys. Rev. Research* **3**, 043014 (2021).

Appendix A: Asymptotic CV-QKD Secret Key-Rates

In the main-text we consider a realistic link-layer for fully trusted QKD networks, in which all point-to-point links utilise a Gaussian-modulated, coherent-state-based CV-QKD protocol and heterodyne detection [56]. By assuming that the nodes are trusted, we can assume that each link utilises a point-to-point QKD protocol. Furthermore, we consider (i) asymptotic secret key-rates in the asymptotic limit of many system exchanges (discounting the additional detrimental effects of parameter estimation or finite key-lengths) and (ii) all imperfections in the trusted communicators' experimental setup are untrusted, i.e. all photons lost along the channel *and* at the receiver are leaked into Eve's environment.

Let us be more precise: Consider a point-to-point fiber channel connecting Alice and Bob who employ the Gaussian modulated coherent-state CV-QKD protocol outlined in Ref. [56]. Let us denote the total channel loss as $\tau = \eta_{\text{eff}}\eta_{\text{ch}}$ where η_{ch} is the usual fiber-channel loss and η_{eff} is the detector efficiency (which we take to be $\eta_{\text{eff}} = 0.7$). Alice and Bob exchange $\bar{n}_T$ photons from the transmitter, but at the receiver there are $\bar{n}_R$ photons at the receiver, where

$$\bar{n}_R := \tau \bar{n}_T + \bar{n}, \quad (\text{A1})$$

such that $\bar{n} := \eta_{\text{eff}}\bar{n}_{\text{bg}} + \bar{n}_{\text{ex}}$ accounts for additional photons due to inefficiencies, background noise and experimental setup noise. In this protocol, the primary contribution of setup noise is due to the local oscillator, which in the local-local oscillator (LLO) protocol contributes both phase and electronic errors (see Ref. [56] for more details).

Using the protocol's entanglement-based representation, and assuming a coherent state modulation of μ , Alice and Bob share the Gaussian state ρ_{AB} after the thermal-loss channel with covariance matrix,

$$\mathbf{V}_{AB} := \begin{pmatrix} \mu \mathbf{I} & \sqrt{\tau(\mu^2 - 1)} \\ \sqrt{\tau(\mu^2 - 1)} & b \mathbf{I} \end{pmatrix} \quad (\text{A2})$$

where $b := \tau(\mu - 1) + 2\bar{n} + 1$. Given that $\nu_{\pm}$ denotes the symplectic eigenvalues of $\mathbf{V}_{AB}$, Eve's Holevo bound is given by

$$\chi_E := h(\nu_+) + h(\nu_-) - h\left(\mu - \frac{\tau(\mu^2 - 1)}{b + 1}\right). \quad (\text{A3})$$

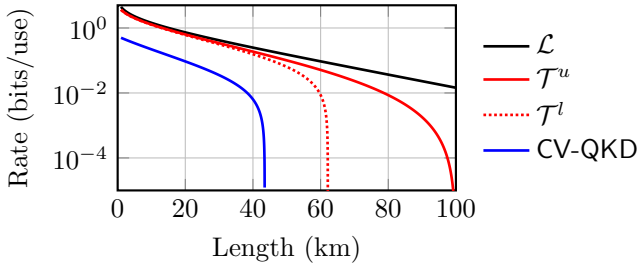


Figure 6. Behaviour of point-to-point bosonic pure-loss capacity $\mathcal{L}$, thermal-loss capacity bounds $\mathcal{T}^{u/l}$ and asymptotic CV-QKD rate with respect to channel length.

where $h(x) := \frac{x+1}{2} \log_2(\frac{x+1}{2}) - \frac{x-1}{2} \log_2(\frac{x-1}{2})$. Consequently, the secret key-rate is

$$K_{\tau, \tilde{n}} := \beta I_{AB}^{\tau, \tilde{n}} - \chi_E^{\tau, \tilde{n}} \quad (\text{A4})$$

where $I_{AB} := \log_2 \left(1 + \frac{\tau(\mu-1)}{2(\tilde{n}+1)} \right)$ is the mutual information between Alice and Bob and $\beta \in [0, 1]$ is the reconciliation efficiency which accounts for imperfect data-processing (which we approximate as $\beta = 0.95$).

This point-to-point rate can then be used to describe the communication rate along a single network edge, $K_{xy} = K_{\tau_{xy}, \tilde{n}_{xy}}$ and provide an accurate link-layer described for trusted node QKD networks. Fig. 6 displays a comparison of the single-edge rates asymptotic CV-QKD rate, bosonic thermal-loss capacity bounds and the exact bosonic pure-loss capacity with respect to channel length.

Appendix B: Network Cuts

An important graph-theoretic concept in the context of network performance is that of *cuts* and *cut-sets*. Consider a network $\mathcal{N} = (P, E)$ with two remote end-users $\mathbf{a}, \mathbf{b} \in P$. We define a cut C as a bipartition of all network nodes P into two disjoint subsets (P_a, P_b) such that the end-user nodes become completely disconnected, $\mathbf{a} \in P_a$ and $\mathbf{b} \in P_b$, where $P_a \cap P_b = \emptyset$. A cut C generates an associated cut-set; a collection of network edges $\tilde{C}$ which when removed cause the partitioning,

$$\tilde{C} = \{(\mathbf{x}, \mathbf{y}) \in E \mid \mathbf{a} \in P_a, \mathbf{b} \in P_b\}. \quad (\text{B1})$$

Under the action of a cut, a network is successfully partitioned

$$\mathcal{N} = (P, E) \xrightarrow{\text{Cut: } C} (P, E \setminus \tilde{C}) = (P_a \cup P_b, E \setminus \tilde{C}), \quad (\text{B2})$$

so that there no longer exists a path between $\mathbf{a}$ and $\mathbf{b}$. Network cuts play a key role in the derivation of end-to-end network rates.

Appendix C: Generalisations of Dijkstra's Algorithm

It is possible to construct general versions of DA in which one optimises path-wise properties associated with

end-to-end routes. It is well known that DA can be used to minimize path length (shortest path problem) or modified to maximize the bottleneck rate (widest path problem). But one can be more general and can define a global cost function $\mathcal{F}_\omega$ such that goal of DA is to find

$$\omega^* = \arg X \mathcal{F}_\omega, \quad (\text{C1})$$

such that $X \in \{\min, \max\}$. In this way, $\mathcal{F}_\omega$ may admit a more complex characterisation of routing value.

Consider an N -node network $\mathcal{N} = (P, E)$. Let us define an N element *tentative path cost set* $T = \{T_x\}_{x \in P}$, which will be used to track the cost of routing throughout the search. This set will be initialised with a unique value for the source node $T_s = \varepsilon_{\text{init}}$, while all other nodes are initialised with a different value $T_x = \varepsilon'_{\text{init}}$ for all $x \in P \setminus \{s\}$. Thus, T takes the initial form

$$T = \{\varepsilon'_{\text{init}}, \dots, \varepsilon'_{\text{init}}, \underbrace{\varepsilon_{\text{init}}}_{\text{source}}, \varepsilon'_{\text{init}}, \dots, \varepsilon'_{\text{init}}\}. \quad (\text{C2})$$

Along with this information, we also have access to point-to-point properties of each edge in the network. There may exist m different types of properties for every edge $(\mathbf{x}, \mathbf{y}) \in E$, which we denote via the set $\{c_{xy}^1, \dots, c_{xy}^m\}$.

The algorithm then operates greedily; starting at an initial source node $\mathbf{s}$, it traverses throughout the network hopping from node to node while keeping track of a tentative path cost from $\mathbf{s}$ to any other network node $\mathbf{x} \in P$. Suppose that we have traversed from $\mathbf{s}$ to a node $\mathbf{x}$ and must decide whether to move to its neighbour $\mathbf{y}$ or not. Here, we must evaluate the tentative cost to the path if we make this move, since we should only hop to a subsequent node if it optimises the routing cost. The tentative path cost is evaluated via a *tentative cost function* F_ω which admits the form

$$F_\omega^{s, x \rightarrow y} = F_\omega(T_x, T_y, \{c_{xy}^1, \dots, c_{xy}^m\}). \quad (\text{C3})$$

That is, F_ω is a function of T_x the tentative path cost from $\mathbf{s} \rightarrow \mathbf{x}$, T_y the tentative path cost from $\mathbf{s} \rightarrow \mathbf{y}$ and the known point-to-point properties associated with moving along the edge $(\mathbf{x}, \mathbf{y}) \in E$. Note that F_ω is not equal to the global cost function $\mathcal{F}_\omega$, but they are inexorably tied. The tentative cost function is simply a translation of the global cost function into a form that can be employed within Dijkstra's algorithm.

With the ability to evaluate the impact this hop has on the total path, the value $F_\omega^{s, x \rightarrow y}$ can be compared with the tentative path cost T_y . If the goal is to minimize the path cost, then the search will hop to $\mathbf{y}$ iff $F_\omega^{s, x \rightarrow y} < T_y$ i.e. moving to this edge better minimizes the cost function. Contrarily, if the goal is to maximize the path cost function then the search will hop to $\mathbf{y}$ iff $F_\omega^{s, x \rightarrow y} > T_y$. The algorithm follows these steps, evaluating, hopping and reevaluating the path cost until eventually the target node $\mathbf{t}$ is reached. Since the algorithm is greedy, it follows an optimal path throughout the network at all times. Once it arrives at a node, one can then be sure that the path followed has been the best one.

Algorithm 1 General Dijkstra Algorithm

Inputs: Network - $\mathcal{N} = (P, E)$, Source - $\mathbf{s}$, Target - $\mathbf{t}$,
 Min/Max Functions - $(X, f_X) \in \{(\min, <), (\max, >)\}$.
 Tentative Cost Function - F_ω ,

```

1: procedure MINCOSTPATH( $\mathbf{s}, \mathbf{t}, \mathcal{N}$ )
2:   Priority queue  $Q = \{\mathbf{x}\}_{\mathbf{x} \in P}$ 
3:   Parent node set:  $P = \{\text{undef}\}_{\mathbf{x} \in P}$ 
4:   Tentative cost set:  $T = \{\varepsilon'_{\text{init}}\}_{\mathbf{x} \in P}$ 
5:    $T_{\mathbf{s}} \leftarrow \varepsilon_{\text{init}}$ 

6:   while  $|Q| > 0$  do
7:      $\mathbf{u} \leftarrow \arg \max_{\mathbf{x} \in Q} T_{\mathbf{x}}$  (and remove  $\mathbf{u}$  from  $Q$ )
8:     if  $\mathbf{u} \neq \mathbf{t}$  then
9:       for all neighbours  $\mathbf{v} \in Q$  of  $\mathbf{u}$  do
10:         $a \leftarrow F_\omega(T_{\mathbf{u}}, \{c_{\mathbf{uv}}^i\}_{i=1}^m)$ 
11:        if  $f_X(a, T_{\mathbf{v}})$  is true then
12:           $T_{\mathbf{v}} \leftarrow a$ ,  $P_{\mathbf{v}} \leftarrow \mathbf{u}$ 
13:          Reprioritise  $Q$  wrt  $T$ 

  return CONSTRUCTPATH( $P, T$ )
  
```

Figure 7. Generalised Dijkstra’s algorithm for end-to-end route optimisation with respect to a cost function $\mathcal{F}_\omega$. Any cost function has a tentative counterpart $F_\omega^{s, \mathbf{x} \rightarrow \mathbf{y}}$ which is used to evaluate movement throughout the network. The above pseudocode describes the network exploration phase which is followed by a CONSTRUCTPATH subroutine which simply back tracks from the target node $\mathbf{t}$ to $\mathbf{s}$ using the constructed tentative cost and parent node sets.

Let us specify to a couple of examples. When minimizing the path length the global and tentative cost functions take the form

$$\begin{aligned} \mathcal{F}_\omega &= \sum_{(\mathbf{x}, \mathbf{y}) \in \omega} l_{\mathbf{xy}}, \\ F_\omega^{s, \mathbf{x} \rightarrow \mathbf{y}} &= T_{\mathbf{x}} + l_{\mathbf{xy}}, \end{aligned} \quad (\text{C4})$$

where $l_{\mathbf{xy}}$ is a measure of point-to-point length of the edge $(\mathbf{x}, \mathbf{y}) \in E$. To minimize this quantity we use the initialisation values are $(\varepsilon_{\text{init}}, \varepsilon'_{\text{init}}) = (0, \infty)$.

For maximizing the bottleneck rate we define the global and tentative cost functions,

$$\begin{aligned} \mathcal{F}_\omega &= \min_{(\mathbf{x}, \mathbf{y}) \in \omega} K_{\mathbf{xy}}, \\ F_\omega^{s, \mathbf{x} \rightarrow \mathbf{y}} &= \max(T_{\mathbf{y}}, \min(T_{\mathbf{x}}, K_{\mathbf{xy}})), \end{aligned} \quad (\text{C5})$$

where $K_{\mathbf{xy}}$ is the point-to-point rate of the edge $(\mathbf{x}, \mathbf{y}) \in E$. Here, we use $(\varepsilon_{\text{init}}, \varepsilon'_{\text{init}}) = (\infty, -\infty)$.

Appendix D: The Multiple Disjoint Paths Algorithm (MDPAlg) and its variant

1. MDPAlg

As discussed in the main text, the MDPAlg offers an efficient means of determining multiple disjoint paths from a source $\mathbf{s}$ to a target $\mathbf{t}$ (or targets $\{\mathbf{t}_i\}_i$) within a network. While DA evaluates and stores the minimum cost to travel between the source and target nodes, the MDPAlg acquires additional information about the *accumulated* cost of traversing from a source to target node through any of its neighbours [39]. This is enormously useful, and provides a mechanism for identifying many cost efficient routes between the source and target.

Much like DA, the algorithm is split into two phases, (i) network exploration and (ii) path reconstruction. Throughout network exploration the algorithm proceeds similarly to DA in which the tentative cost matrix T is constructed. The primary difference between DA and the MDPAlg is found on line 11 of Fig. 7: Even when the computed tentative cost a is not considered optimal, in the MDPAlg it is stored as an off diagonal element of $T_{\mathbf{uv}}$, describing the accumulated cost of travelling from node $\mathbf{s}$ to $\mathbf{v}$ via $\mathbf{u}$.

This additional information is then used in the path reconstruction phase to identify many end-to-end routes. The reconstruction is sequential and straightforward: Starting at a target node $\mathbf{t}$ a path is built by moving to the neighbour which incurs the lowest cost in the tentative cost matrix, T , until the source node is reached. On the first path reconstruction, this is simple and the minimum cost path ω_1 is produced, which adds edges to the routing edge set ($E_\omega = \omega_1$). In order to construct subsequent disjoint paths, one must then enforce disjointedness by restricting the use of any edges which were used previously (for link disjointedness). In other words, future reconstructions will ignore the cost matrix elements $T_{\mathbf{xy}}$ for $(\mathbf{x}, \mathbf{y}) \in E_\omega$. Repeated this process, path reconstruction may occur by moving from the target $\mathbf{t}$ to its *next* minimum cost neighbour, and so on until $\mathbf{s}$ is met again.

Through this sequential process of path building coupled with edge restriction, a number of edge-disjoint paths can be built which have favourable properties. The node-disjoint variant of the MDPAlg can be easily achieved through a small modification; instead of only restricting previously used edges from subsequent path reconstructions, one must restrict the use of any edges connected to nodes used in the previous routes.

2. Rate Maximisation and IAR-MDPAlg

The MDPAlg was introduced with the intention of identifying multiple disjoint *shortest* paths which minimize the cumulative cost of edge-weights along an end-to-end route. Consequently, it is not immediately clear

how the MDPAlg can be translated for the purposes of rate maximisation. One might assume that since there exists a variant of DA for this purpose (the widest path algorithm outlined in Appendix C) then it should be able to modify the MDPAlg in an identical way.

Unfortunately, this is not the case and we must be careful in building this variant (IAR-MDPAlg). The widest path version of DA locates an end-to-end route which maximizes a bottleneck rate between source and target nodes. It does so using via the generalised DA and the cost functions in Eq. (C5).

Now let us consider the scenario in which we possess a source node s , target node t and wish to find multiple end-to-end routes which maximize their bottleneck rates, e.g. in the vein of iterative Dijkstra. We want to employ the more efficient MDPAlg to do this, in which only a single Dijkstra search is necessary, followed by path reconstruction. In this context, what information would a tentative cost matrix, T , collect? The diagonal element T_{xx} would contain the minimum cost associated with routing between s and x , i.e. the bottleneck rate along the optimal route. Meanwhile, the off-diagonal quantities T_{xy} indicate the potentially sub-optimal cost incurred when routing between x and s through the intermediate node y .

On the first iteration of path reconstruction, it is easy to identify the minimum cost path ω_1 by backtracking a route to its most favourable neighbour node n_1 and repeating this process according to subsequent nodes and the diagonal elements of T . Since we wish to construct edge-disjoint paths, any edge contained within ω_1 is then rendered inaccessible by future routes. Following ω_1 , we wish to construct a second end-to-end path using information from the cost matrix. This is initiated by hopping from the target t to its next more favourable neighbour, which we label n_2 . The cost matrix has collected seemingly valuable information about how to best traverse between s and t through n_2 .

This is where our issue begins to emerge: What if the best path between n_2 and s shares edges with ω_1 (the first minimum cost path)? Widest paths are often highly degenerate (especially in large networks) due to the fact that their value is characterised completely by single-edge rates. The IAR-MDPAlg will “push” the second path towards the optimal path (where the edges are shared) but must then move along alternative edges due to the edge-disjointness restriction, i.e. it is then “pulled” away from the optimal path. These alternative edges may possess poor rates, but the algorithm will nonetheless use them because of its now unstructured method of path building, and the cost matrix does not contain enough information to discourage their use. The push-and-pull continues until s is reached, by which point the route has been compromised and possesses a poor bottleneck rate. This push-and-pull effect is only avoided when the optimal end-to-end route from t to s through a neighbouring node is edge-disjoint with the optimal route, which is rare in large-scale, highly connected networks.

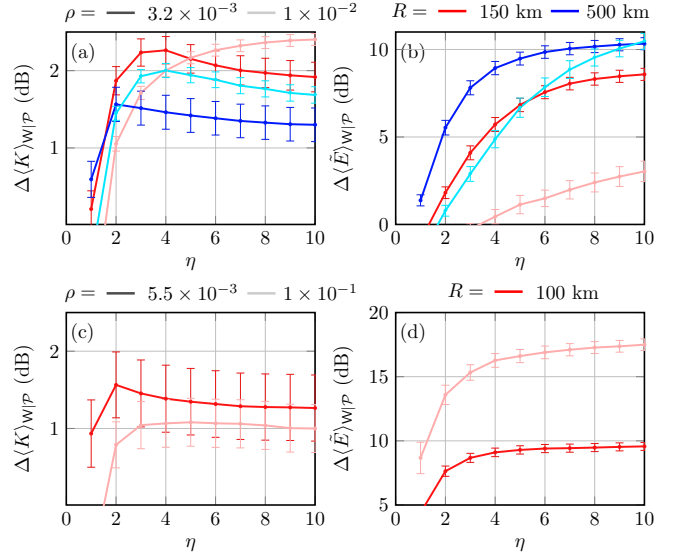


Figure 8. Behaviour of IAR ansatz with respect to inverse rate penalty η and fixed edge-usage penalty $\epsilon = 1$ over the classes of bosonic thermal-loss (a)-(b) Waxman networks and (c)-(d) scale-free networks ($\sigma_r = 1$) using the thermal single-edge capacity upper-bounds. For a number of nodal densities ρ and network radii, Panels (a),(c) plots the ensemble average rate amplification $\Delta\langle K\rangle_{N|P}$ defined in Eq. (D3) gathered by the $M = 2$ MDPAlg protocol using the inverse rate-sum ansatz, with respect to variable η . Panels (b),(d) plots the ensemble average consumption amplification $\Delta\langle \tilde{E}\rangle_{N|P}$ defined in Eq. (D4) with respect to η .

As a result, the bottleneck rate cost functions used within DA are not so effective in the present context. One should avoid the employment of routing cost functions that possess high degeneracies, such as the direct translation of the widest path cost functions. Instead, one should explore cumulative costs, similar to that used in the shortest path formulation. Cumulative cost functions which capture properties of entire paths are more suited to this algorithm and can more effectively motivate end-to-end routing.

3. Inverse-Accumulated-Rate Ansatz

To achieve rate maximisation, a potential ansatz for the tentative cost function may take the form,

$$F_{\omega}^{s,x \rightarrow y} = T_{xx} + K_{xy}^{-\eta} + \epsilon. \quad (D1)$$

which corresponds to minimizing the following global cost function,

$$\mathcal{F}_{\omega} = \sum_{(x,y) \in \omega} (K_{xy}^{-\eta} + \epsilon). \quad (D2)$$

This aligns with minimizing the sum of the inverse point-to-point rates along a path ω , adding a penalty for channel usage ϵ associated with each link. Minimizing this

cost function simultaneously locates a path which maximizes the sum of the point-to-point rates along the path while minimizing its length; *indirectly* identifying a high-rate and edge-efficient route.

In this tentative cost function, $\eta, \epsilon \in \mathbb{R}_0^+$ are hyperparameters used to impose a tradeoff between rate and path length. The parameter η can be thought of as a rate motivator; if η is large, then the inverse term $K_{xy}^{-\eta}$ forces a large penalty to the cost function when low rate edges are included in the route. If the single edge rate is large, this incurs a low cost, motivating the use of an edge. Meanwhile, ϵ enforces a constant penalty term for edge usage, thus encouraging efficient routing. The edge-usage penalty can be generalised as an edge-wise property ϵ_{xy} such as spatial length etc., to more accurately manage routing efficiency. The ability to control this tradeoff is extremely useful and is typically ignored in standard rate-maximisation algorithms.

With some simple numerical inspection, we see that the inverse rate-sum ansatz can achieve strong rate and resource management. In Fig. 8 we perform an analysis of the ansatz given in Eq. (D1) for both Waxman networks and scale-free models. Here, we explore the minimum improvement offered via an MDPAlg protocol with two fixed routes, $\mathcal{P}_{\text{mdp}}^{M=2}$, by measuring the *amplification* of its ensemble average rate and routing consumption over the optimal single-path protocol (measured in decibels),

$$\Delta\langle K \rangle_{N|\mathcal{P}} := 10 \log_{10} \left[\langle K \rangle_{N|\mathcal{P}_{\text{mdp}}^{M=2}} / \langle K \rangle_{N|\mathcal{P}_{\text{sp}}} \right], \quad (\text{D3})$$

$$\Delta\langle \tilde{E} \rangle_{N|\mathcal{P}} := 10 \log_{10} \left[\langle \tilde{E} \rangle_{N|\mathcal{P}_{\text{mdp}}^{M=2}} / \langle \tilde{E} \rangle_{N|\mathcal{P}_{\text{sp}}} \right]. \quad (\text{D4})$$

When $\Delta\langle K \rangle_{N|\mathcal{P}} > 0$ then the MDPAlg protocol offers a rate advantage over the optimal single-path protocol. However, the magnitude of $\Delta\langle \tilde{E} \rangle_{N|\mathcal{P}} > 0$ quantifies the *increase* in routing consumption required to achieve this rate advantage. Optimally tuning the inverse rate-sum ansatz corresponds to maximizing $\Delta\langle K \rangle_{N|\mathcal{P}}$ while minimizing $\Delta\langle \tilde{E} \rangle_{N|\mathcal{P}}$.

Fig. 8 illustrates the behaviour of these quantities with

respect to a variable inverse-rate penalty η , while keeping $\epsilon = 1$ fixed and constant. For both classes of network, it is clear that for the scenarios considered in this work that the inverse rate-sum ansatz can be used to effectively enhance performance over that of single-path algorithms without incurring considerable routing cost. The algorithm can locate another end-to-end rate that effective enhances the rate while keeping the routing consumption increase at a manageable level. The more η is increased, the greater emphasis it places on maximizing the end-to-end rate; however, this does not always correlate with *obtaining* a greater end-to-end rate. Indeed, for both network models the rate enhancement tends to plateau while the routing consumption may continue to increase leading to wasted network usage. Hence, one should be careful not to increase η too significantly with respect to the edge-usage penalty, ϵ .

Clearly, the optimal hyperparameters depend on the network model, connectivity and nodal density. An ideal implementation of the inverse rate-sum ansatz would fine-tune η and ϵ corresponding to these properties. For simplicity, throughout this work we have employed the quantities $\eta = 5$ and $\epsilon = 1$, motivated by the analyses above as a reasonable, multi-purpose choice with respect to different network characteristics. Future studies should focus on optimizing/controlling these hyperparameters with respect to specific network models or routing strategies.

This is by no means an optimal approach. Nonetheless, the enormous benefit associated with quickly locating additional routes proves to outweigh any weakness in the approximation. In future investigations, it may be interesting to explore more sophisticated cost analyses, e.g. a neural network variational ansatz. Such ansatzes may be of significant benefit when one wishes to optimise more than just rate and path length, e.g. balancing the routing priorities of multiple user pairs, or considering waiting-times in quantum memories for entanglement distribution.