

Article

Topology-Oblivious Random-Walk Key Relaying in Quantum Key Distribution Networks

Krišjānis  Petručeņa ^{*} , Sergejs Kozlovičs , Juris Viksna , Elīna Kalniņa , Reinis Isaks , Edgars Celms ,
Lelde Lāce  and Edgars Rencis 

Institute of Mathematics and Computer Science, University of Latvia, LV-1459 Riga, Latvia; sergejs.kozlovics@lumii.lv (S.K.); juris.viksna@lumii.lv (J.V.); reinis.isaks@lumii.lv (R.I.); edgars.celms@lumii.lv (E.C.); lelde.lace@lumii.lv (L.L.); edgars.rencis@lumii.lv (E.R.)

^{*} Correspondence: krisjanis.petrucena@lumii.lv

Abstract

Quantum key distribution (QKD) networks require relaying when distant key management entities share no direct quantum link. Most relay strategies, however, rely on centralized control or globally maintained routing state. This paper asks whether useful security and efficiency can still be obtained with topology-oblivious stochastic forwarding. It studies the security-overhead trade-off in a model in which fragmented key material is relayed via random-walk variants and reconstructed under privacy amplification. The analysis asks whether strictly local forwarding can retain useful information-theoretic security (ITS). Evaluation on the GÉANT topology, representing a European academic backbone network, shows clear differences between random-walk variants. The proposed highest-score-neighbor local path-diversification heuristic reduces the probability that relayed key material passes through a compromised node. The evaluation also shows that scouting-based loop erasure significantly shortens sampled routes and improves throughput in the model. Against one- to three-node cartels, random flow protects slightly more source–target pairs than a static disjoint-multipath method on the evaluated topologies. These findings position topology-oblivious stochastic forwarding as a simpler decentralized design for QKD relaying without centralized orchestration or gossip protocols.

Keywords: quantum key distribution; QKD networks; trusted-node relaying; random walks; topology-oblivious routing; privacy amplification



Academic Editors: Rosario Lo Franco,
Avishy Carmi and Eliahu Cohen

Received: 1 April 2026

Revised: 10 June 2026

Accepted: 13 June 2026

Published: 16 June 2026

Copyright: © 2026 by the authors.

Licensee MDPI, Basel, Switzerland.

This article is an open access article distributed under the terms and conditions of the [Creative Commons Attribution \(CC BY\)](https://creativecommons.org/licenses/by/4.0/) license.

1. Introduction

Quantum key distribution (QKD) can provide symmetric keys with *information-theoretic security* (ITS) [1], relying on the principles of quantum mechanics, in particular the no-cloning theorem and the disturbance caused by measuring non-orthogonal quantum states. Since the pioneering BB84 protocol [2,3], QKD has evolved from laboratory demonstrations into metropolitan-scale and intercity deployments capable of supporting secure communication infrastructures [4–6]. QKD systems available on the market have limited working distance and secret key rate [7]. Typically, fiber-based QKD links operate at distances smaller than 150 km [8], with some deployments reaching 250 km [9]. To address distance limitations, deployments are organized as multi-hop QKD networks with trusted nodes [10] and a key management layer. Adjacent nodes establish pairwise secret keys using QKD links, while the key management layer is responsible for authorization, key forwarding, and key delivery. Ultimately, any two QKD network nodes—key management

entities (KMEs)—can establish a shared secret key. If we believe that no node is compromised and we use one-time pad (OTP) encryption for hop-by-hop relaying, we can reclaim ITS. In practice, we must consider a threat model where some nodes are compromised and may eavesdrop on secrets. We will refer to a cooperating, malicious node set as “cartel”. In this paper, we ask whether useful *efficiency* (security-overhead trade-off) can be achieved when forwarding is topology-oblivious. If viable, this approach could simplify the decentralized key-management layer, especially in dynamic networks with nodes and edges introduced over time. By contrast, most QKD-network protocol proposals either rely on centralized or SDN-assisted orchestration [11,12], or on topology-aware path selection. In well-known distributed link-state routing protocols such as OSPF [13], routers flood link-state advertisements (LSAs), spreading local link information in a gossip-like manner so that all routers can learn the topology. Afterwards, nodes can compute end-to-end paths using the learned topology. QKD-specific OSPF-style relaying has also been proposed as a distributed alternative to centralized orchestration [14]. In this paper, we propose an alternative random-walk construction and avoid gossip protocols. We study a topology-oblivious stochastic relay scheme that we call a *random flow*. In one random flow from source s to destination t , the source splits key material into M fragments and emits those fragments in parallel. Each fragment is then forwarded independently using a random-walk *variant*. The destination reconstructs the recovered material and applies seeded privacy amplification to derive the final key *block* $\mathcal{K}$ of size θ . In this setting, exposure is the probability that one fragment from s to t traverses at least one relay in a compromised cartel. For an admissible cartel class, $\chi_{s,t}$ denotes the worst such hit probability. We also track expected hop count $h_{s,t}$ and design-threshold efficiency $\eta_{s,t}^\tau$, where τ is a configured upper bound on admissible-cartel exposure and $\eta_{s,t}^\tau$ is a model-based extracted-output proxy per consumed QKD-derived link-key bit. Our contributions are the following:

1. A topology-oblivious random-flow relaying model in which fragmented key material is forwarded by stochastic rules that avoid global adjacency knowledge, link-state maintenance, and end-to-end path computation.
2. A highest-score path-diversification heuristic that improves worst-case exposure without requiring global topology knowledge or even the node count.
3. A scouting-based loop-erasure mechanism that shortens realized payload routes, reduces queueing pressure, and eliminates self-induced cyclic waiting in the model.
4. A simulation study on reconstructed and synthetic topologies that compares the random-walk variants under exposure, hop-count, efficiency, and throughput proxies.

We parameterize the strongest possible topology-independent target by the pair connectivity $\kappa(s, t)$. A cartel of size $\kappa(s, t)$ positioned on a minimum s – t vertex cut can eavesdrop all transmitted key fragments, so the largest universal cartel-size target is $|\mathcal{C}| \leq \kappa(s, t) - 1$. For random flow, the resulting ITS claim is conditional on the measured or assumed exposure of the chosen admissible cartel class being at most the selected threshold τ . In the evaluation, we use $\tau = 98\%$ as the default design threshold, report single-relay exposure for all random-walk variants, and separately give GÉANT sensitivity measurements for non-cut two- and three-node cartels. Those cartel measurements are descriptive; they do not turn the single-relay results into a blanket multi-relay guarantee.

Scope and Limitations

The study stays at the trusted-node key-management layer and does not model physical-layer calibration, synchronization, or alternative physical-layer relaying protocols. GÉANT and NSFNET are structural proxies without inserted intermediate relays for long fiber spans. The throughput model uses a uniform 1 kbit/s secret-key rate and omits classical control overhead. These choices support comparison of local forwarding heuristics,

not a deployment-ready Q-KMS product. The remainder of the paper is organized as follows. Section 2 reviews the QKD-network setting and related work. Section 3 introduces the random-flow model, random-walk variants, privacy-amplification view, and loop-erasure mechanism. Section 4 presents the simulation study and comparative results. Section 5 concludes the paper.

2. Background and Related Work

For QKD network modeling, QKD devices are commonly abstracted as black boxes that provide adjacent nodes with a stream of QKD-link-derived symmetric keys. If one endpoint obtains such a link key, the corresponding endpoint is assumed to have established the same key, and the key is secure according to the underlying QKD protocol, up to its stated security parameters. The BB84 protocol [2,3] is the classic paradigm for quantum key distribution (QKD). Here, classical bits are represented as quantum states, e.g., polarized photons, transmitted over an optical channel. After a process of sifting and reconciliation, the communicating parties can detect the presence of any eavesdropper. One of the principal challenges in terrestrial QKD deployment is the rapid decrease in secret-key rate with increased transmission distance, primarily due to fiber attenuation and other implementation losses. As a result, except in certain experimental satellite-based scenarios [15], practical direct-transmission QKD links are generally restricted to distances of about 150 km [8]. This reflects exponential fiber loss and repeaterless rate–loss bounds for optical channels [16]. Twin-field QKD can exceed the conventional direct-transmission scaling by using single-photon interference at an intermediate measurement node, but current systems remain experimental prototypes [17,18].

2.1. Trusted-Node Relaying

Because of this range limit, larger deployments use relay nodes that pass key material from one hop to the next until it reaches the destination. Taken together, the QKD links and relay-capable nodes form a graph-like QKD network. The tasks of relaying, buffering, and authorizing access to key material are handled by the key-management layer. At that transport layer [19], two nodes A, B are adjacent when they share a direct QKD link. Such neighbors continuously obtain a common sequence of link keys $\{\ell_{A,B}\}$, usually identified by UUIDs. These keys may be buffered and then consumed for OTP encryption of data sent over an authenticated classical channel. To establish an end-to-end key between non-adjacent nodes s and t , the source can generate a fresh final key K_f , choose a relay path, and send $K_f \oplus \ell$ to its first-hop neighbor. Each relay on the path decrypts, recovers K_f , and re-encrypts it for the next hop. This explains why the security model of trusted-node QKD networks depends critically on intermediate relays remaining uncompromised [10].

2.2. ETSI 014 and Q-KMS

The software stack that performs relaying, buffering, and authorization is often referred to as a Quantum KMS (Q-KMS). This differs from a conventional KMS, whose primary role is usually key lifecycle management: key generation, access control, auditing, and cryptographic service exposure, frequently via hardware security modules. In practice, vendors such as ID Quantique, Toshiba, and LuxQuanta typically ship proprietary Q-KMS software stacks alongside their hardware, which can limit interoperability across different vendors. At the application interface, ETSI GS QKD 014 [20] defines a RESTful HTTPS API through which applications obtain QKD-derived keys as JSON containers containing pairs of (UUID, key). The standard separates Secure Application Entities (SAEs) from Key Management Entities (KMEs), with the KMEs placed inside Trusted Nodes together with one or more QKD Entities (QKDEs) that terminate the physical QKD links. Mutual

certificate-based authentication between SAE and KME is part of the model, but inter-KME relay protocols are left unspecified. Figure 1 sketches this separation between the standardized application-facing API and the non-standard relay logic behind it.

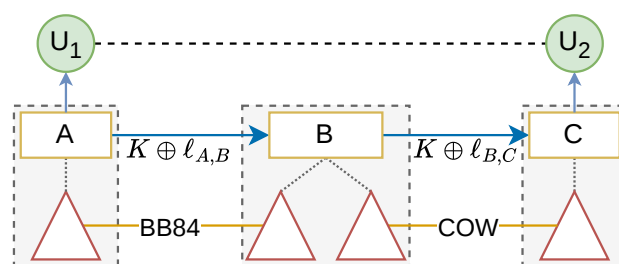


Figure 1. Illustration of ETSI 014-style key delivery. User U_1 requests a fresh key from KME A and passes the returned key identifier to user U_2 , which later retrieves the same key from KME C. Letters A–C label the three KMEs/trusted nodes in the example relay chain. The corresponding KMEs are embedded in trusted nodes with QKD link endpoints, while the underlying Q-KMS transfers the final key K_f hop-by-hop using OTP protection with $\ell_{A,B}$ and $\ell_{B,C}$.

This interface should be understood as a source of key material, not as a complete substitute for end-to-end application or network security. A Secure Application Entity should not rely solely on the QKD network layer: for defense in depth, the communicating applications should still run a fresh post-quantum or hybrid key exchange in the protecting protocol, for example in a TLS 1.3 handshake or IPsec/IKEv2 exchange, and combine that result with QKD-derived material when available [21–23].

2.3. Multipath Routing

To reduce dependence on any single intermediate “trusted” node, the end-to-end key K_f may be decomposed and sent over several strictly or partially node-disjoint routes. This family of approaches is usually called multi-path QKD. In the simplest case, the destination reconstructs K_f from independently delivered pieces, for example by combining $K_1, \dots, K_n$ as $K_f \leftarrow K_1 \oplus \dots \oplus K_n$. The adversary then has to compromise at least one relay on every relevant path to recover the final key. More general (t, n) secret-sharing constructions offer tolerance against partial share leakage, albeit with additional overhead [24,25]. Such techniques are a standard way to reason about partially trusted QKD networks. Among topology-aware alternatives, deterministic routing over multiple node-disjoint paths is the main benchmark for our stochastic relaying model [26]. Figure 2 illustrates the connectivity condition. For endpoints s and t with $\kappa(s, t) = 3$, the source can XOR-split K_f across three internally node-disjoint routes. The adversary must then compromise at least one relay on every path to recover K_f . Redundant parallel links between s and t do not by themselves increase the number of node-disjoint routes. A cartel of size $m = 2$ cannot meet all three routes, whereas a cartel of size $m = \kappa(s, t)$ placed on a minimum s – t vertex cut intercepts every one of them.

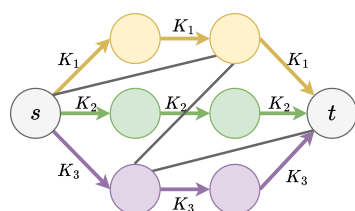


Figure 2. Three internally node-disjoint s – t paths (colored) with $\kappa(s, t) = 3$. The colored curves distinguish the three separate paths taken by K_1 , K_2 , and K_3 . A cartel of size 2 cannot eavesdrop on all key shares, but a cartel of size 3 on a minimum vertex cut can.

For a given pair (s, t) , minimum-total-length node-disjoint path sets are computable in polynomial time by Suurballe's algorithm and Bhandari's node-disjoint variant [27,28], and we adopt that construction as the multipath baseline in Section 4.7. Most routing proposals for trusted-relay QKD networks are topology-aware and assume some form of global or near-global state, often maintained through link-state-style control exchange and constrained by consumable key material, trust exposure, dynamic link availability, and multipath overhead [12,29,30]. Recent work has also proposed an explicitly OSPF-based distributed key-relay architecture for QKD networks, directly contrasting with our aim to avoid a link-state control plane [14]. Recent pre-relaying and decentralized key-pool work studies a different resource-management direction: QKD key material is buffered, non-reusable, and scarce, so proactive virtual key pools can reduce delay or improve allocation success at the cost of storage and relaying overhead [31,32]. Partially trusted approaches instead split key material across multiple paths or shares [24–26,33], while stochastic routing has also been explored before [34,35]. Recent zero-trust proposals pursue a stronger goal than the one studied here [36].

2.4. MDI-QKD

Measurement-device-independent QKD (MDI-QKD) offers a contrasting physical-layer design, building on side-channel-free entanglement-swapping ideas [37] and the practical protocol of Lo, Curty, and Qi [38]. In a star-style architecture, end users operate low-cost transmitters while a shared relay performs Bell-state measurements on arriving pulses [38]. The measurement relay may be untrusted: detector-side-channel attacks, including blinding and click tailoring, do not break confidentiality under the protocol model [38]. This removes trust in the detection hardware at the hub, but not in sources, channels, or classical post-processing. The protocol was first demonstrated experimentally in 2013 [39]. Compared with OTP relaying over trusted nodes, pure MDI-QKD trades a different security profile for demanding synchronization and generally lower key rates and network capacity [40,41]. Neither MDI nor trusted relaying alone achieves arbitrary long-distance scalability with current technology. Recent network studies therefore treat MDI as trusted-node *reduction* rather than elimination, and propose hybrid trusted/untrusted deployments to balance security with key-provisioning reach [42]. Operational models are converging on partially trusted relay architectures. In such networks, MDI or otherwise untrusted segments coexist with conventional trusted OTP relays, and routing selects among them according to distance, key-pool state, and residual trust [33,43]. Attacks can also exploit imperfect relay security at the routing layer even when individual QKD links appear sound [44]. The present work stays at the key-management layer for trusted-node OTP relaying and studies semi-trusted resilience via multipath diversification and topology-oblivious forwarding; it is complementary to, not subsumed by, MDI-QKD access designs.

3. Random Flow

At a high level, we replace deterministic routing by stochastic forwarding of fragmented key material: each of the M fragments is routed by a random walk *variant* from source s towards target t . Node t collects all fragments and derives the final key *block* $\{K_i\}$ of size θ via seeded privacy amplification. We find that the random walk (RW) *variant* plays a significant role in *exposure* $\chi_{s,t}$ —the probability of traversing some malicious node, if cartel is positioned optimally for the given (s, t) pair. Building on the QKD transport model from Section 2, we now switch to the algorithmic abstraction. We represent the QKD network as a simple, undirected graph $G = (V, E)$, where V denotes the set of trusted nodes and E denotes the set of QKD links. The per-link keys $\ell_{u,v}$ introduced earlier are treated here simply as the hop-by-hop OTP resource available on edge (u, v) . To simplify

the notation in the analysis, we assume that both the link-key blocks and the derived end-to-end key have lengths of 256 bits. We assume a reactive (rather than a proactive) model. In reactive mode, key *allocation* request triggers a *transmission*, and transmissions are not initiated ahead of time. A transmission (node $s \rightarrow$ node t) is a key-relaying session by which the final key *block* $\mathcal{K}$ of size θ is established. In the base model (without loop erasure), the source s emits M raw fragments $f_1, \dots, f_M$, which eventually reach destination t by means of random walks. Later (Sections 3.4 and 4.2), we adopt $\tau = 98\%$ as the default design threshold (Table 1). At $M = 1024$, the binomial lower bound is $g_{99.99\%}^*(1024, 98\%) = 6$, which after Theorem 1's privacy-amplification penalty (with $\varepsilon_{\text{PA}} = 2^{-256}$) corresponds to $\theta = g^* - 2 = 4$ extractable 256-bit keys per block. Consequently, the ITS-oriented evaluation at $\tau = 98\%$ uses $M = 1024$; smaller fragment counts are only meaningful for this theorem when paired with a lower validated exposure threshold.

Table 1. Maximum g values such that $P[\# \geq g] > \alpha$ holds for various exposure values χ . These values determine the guaranteed safe-fragment count used by the ITS parameter choice.

M	$g_{99\%}^*(M, \chi)$						$g_{99.99\%}^*(M, \chi)$					
	98%	95%	90%	85%	80%	75%	98%	95%	90%	85%	80%	75%
32	0	0	0	1	2	3	0	0	0	0	0	0
64	0	0	2	4	6	8	0	0	0	1	3	5
128	0	1	6	10	16	21	0	0	2	6	10	15
256	1	5	15	26	37	48	0	2	10	19	29	39
512	4	15	36	59	82	106	1	9	28	48	70	93
1024	11	36	81	128	175	224	6	27	69	113	159	206

3.1. Threat Model and Objectives

For information-theoretic security (confidentiality), the adversary controls a cartel $C \subseteq V \setminus \{s, t\}$ of compromised relays. Compromise reveals each relay's full internal state and the plaintext of every fragment that traverses any relay in C . A fragment is treated as compromised once its realized walk intersects C . Public-key wrapping of fragments does not satisfy the information-theoretic goal of QKD and is not considered in this paper. We assume honest endpoints (s and t), authenticated classical channels based on PKI and mutual TLS, and trusted QKD hardware, excluding side-channel attacks [45]. Theorem 1 should therefore be read as a confidentiality theorem under authenticated transport and successful protocol completion, not as an availability guarantee against Byzantine relays. Let $\kappa(s, t)$ denote the maximum number of internally node-disjoint s – t paths, equivalently the minimum size of an s – t vertex cut. The universal ITS target for a pair (s, t) is protection against every relay cartel of size

$$|C| \leq \kappa(s, t) - 1.$$

This is the strongest possible topology-independent cartel-size target for that pair: a cartel of size $\kappa(s, t)$ placed on a minimum s – t cut can intercept every s – t route. For cartels larger than $\kappa(s, t) - 1$, we do not claim a universal ITS guarantee. The design parameter for random flow is the exposure threshold τ . For an admissible cartel class $\mathcal{C}_{s,t}$ and a realistic topology family, the intended parameter choice must satisfy

$$\max_{C \in \mathcal{C}_{s,t}} p_C^{s,t} \leq \tau,$$

where $p_C^{s,t}$ is the probability that one random-flow fragment intersects C . The value τ is therefore not a universal constant of the protocol. It must be chosen high enough for the evaluated topology and cartel class, and is then used in the binomial lower bound $g_\alpha^*(M, \tau)$

and in the privacy-amplification output length. If the measured worst-case cartel exposure on a topology is above 98%, then $\tau = 98\%$ is not a valid ITS design point for that topology and cartel class. In particular, achieving the $\kappa(s, t) - 1$ target on GÉANT may require experiments at $\tau > 98\%$. In parallel, we also evaluate a placement-specific non-cut cartel class for $|C| \leq 3$. Here, C is admissible for a given (s, t) only when s and t remain connected in $G - C$. This non-cut condition means that at least one uncompromised s – t path still exists, even when a particular non-cut placement has size above the universal all-placements target $\kappa(s, t) - 1$. For these non-cut placements, the ITS claim is again conditional on the measured cartel exposure satisfying $p_C^{s,t} \leq \tau$ and on the fragment independence model used by $g_\alpha^*(M, \tau)$. Choosing τ below the true exposure invalidates the ITS parameter claim because it overstates the lower bound on full-entropy fragments. It does not automatically imply total key disclosure. If at least one full-entropy fragment avoids the cartel, the construction retains the strictly non-ITS computational fallback described in Section 3.5. This fallback is a residual-confidentiality statement based on a preimage-resistant hash or KDF over the received fragment string, not an information-theoretic guarantee. A compromised relay may also drop, delay, misroute, or inject traffic. We treat such active behavior as an availability and integrity issue, not as part of the ITS confidentiality proof. Drop and delay are availability problems and may cause the transmission to fail to complete. Misrouting is observable through the route metadata reported by the destination and can be flagged at the application layer. Injected or modified traffic is rejected by authenticated classical-channel mechanisms, for example per-hop MACs or AEAD tags derived from link material, together with end-to-end route validation where applicable. The expected outcome of any such active misbehavior, when detected, is rejection or abortion of the transmission rather than silent confidentiality loss. Equivalently, an attacker who can disable QKD links should only be able to reduce availability for a given $s \rightarrow t$ transmission, not improve its confidentiality advantage beyond the configured exposure model. The design goals are:

1. Information-theoretic security for admissible cartels whose exposure is bounded by τ .
2. Exposure upper bound τ calibrated for realistic medium-scale QKD networks.
3. No per-transmission advantage for attackers who can disable relay links and nodes.
4. Independent stochastic forwarding that utilizes partial disjointness for security.
5. Compatibility with ETSI GS QKD 014 key delivery API.

Forwarding decisions use neighbor information and variant-specific token state; the node-coloring variant introduced below additionally assumes that the source knows the global node-identifier universe. The topology itself may change over time. We prove the ITS statement for the admissible cartel class through the number of realized fragments that avoid the worst admissible cartel. Most realistic mid-term QKD networks are expected to remain sparse because of the high cost of QKD equipment, deployment, and operation. As a result, pairwise connectivity is often low, and the achievable cartel tolerance is inherently limited by min-cut constraints. For biconnected pairs, $\kappa(s, t) \geq 2$, so the universal target always includes at least one compromised relay and grows with the pair's actual vertex connectivity. Current QKD relay nodes are typically deployed as “trusted nodes” in physically protected facilities, such as telecom points of presence, or secured data centers. Their attack surface is narrower, physical access is controlled, and operational procedures are usually stricter. Consequently, compromise is more likely to arise from a configuration error, insider misuse, supply-chain weakness, or a single negligent administrator. From an operational perspective, multiple simultaneous compromises of relay nodes would represent a high-impact security incident that network operators are strongly motivated to detect and contain quickly. Since QKD networks are expensive, specialized infrastructures with relatively few nodes, they are likely to receive closer monitoring and tighter administrative control than commodity networks. If malicious nodes are not placed optimally on

a minimum cut [46] for a given source–target pair, a larger non-cut cartel may still leave enough unseen fragments for ITS under the same τ -bounded exposure condition.

3.2. Random Walk Notation

We represent each in-flight (travelling from s to t) fragment f_i by a *token*. Concretely, transmission token i is a tuple consisting of id —transmission identifier, source and target pair (s, t) , i —token index, f_i —fragment itself, and σ_i —variant-specific token-local state. State σ_i is the only element that may change over time.

$$\text{token}_i = (\text{id}, s, t, i, f_i, \sigma_i),$$

Each token i is forwarded by a discrete-time walk. The trajectory of token i during the sampled walk is $X_0, X_1, \dots, X_h$, with $X_0 = s$, $X_{k+1} \in N(X_k)$, $X_h = t$, where $N(X_k)$ is the neighborhood of X_k . The whole trajectory is a sampled random sequence $(X_k)_{k=0}^h$, where h is the hop count. The probability distribution from which X_{k+1} is sampled is P_{variant} .

$$X_{k+1} \sim P_{\text{variant}}(\cdot \mid X_k, \sigma_k),$$

It depends only on the current token i and variant-specific state σ . Relays are therefore stateless under this objective and under our interpretation of the random-walk setting. Different fragments use independent forwarding randomness, including independent diversification seeds for HS, so that the exposure events used by the extraction analysis are modeled as independent. The dot in $P_{\text{variant}}(\cdot \mid \dots)$ is a placeholder, and altogether the expression means “distribution over possible next nodes”. The walk terminates when it first hits the target. If the target t is an adjacent neighbor of the current relay v , i.e., $t \in N(v)$, the relay forwards directly to t instead of sampling another stochastic hop. This direct-target rule is the model used in Section 4. For readability, we will adopt the notation $P(v \rightarrow u \mid \star)$ instead of $P(X_{k+1} \mid X_k, \sigma_k)$, where $v = X_k$, $u = X_{k+1}$ and $\star$ is syntactically substituted for some predicate that can be evaluated from σ_k . It is the probability to go from v to u given position X_k and state σ_k .

3.3. Base Random Walk Variants

Simple random walk (R). The simple random walk variant R is memoryless: at node v , the token chooses the next hop uniformly at random.

$$P_R(v \rightarrow u) = \begin{cases} 1/|N(v)|, & u \in N(v), \\ 0, & \text{otherwise.} \end{cases}$$

Non-backtracking random walk (NB). Non-backtracking NB suppresses immediate return. The token state σ carries a single field $\text{prev} \in V \cup \{\text{null}\}$ with the previous node (or null at the start). Let $N'(v) = N(v) \setminus \{p\}$, where $p = \text{prev}$. Then,

$$P_{\text{NB}}(v \rightarrow u \mid p) = \begin{cases} 1/|N'(v)|, & u \in N'(v), \\ 1, & u = p \text{ and } d(v) = 1, \\ 0, & \text{otherwise.} \end{cases}$$

After choosing u , the token updates $\text{prev} \leftarrow v$. On regular expander graphs, NB can “mix” provably faster than R [47]. Informally, an *expander* is a sparse graph with strong connectivity.

Least-recently-visited walk (LRV). LRV biases the walk away from recently visited vertices. We use an LRV-vertex rule: token i maintains timestamps $\text{last}: V \rightarrow \mathbb{N}_0$, where $\text{last}[x]$ is the most recent time at which the token visited x . We initialize $\text{last}[s] = 1$, and return $\text{last}[x] = 0$ by default. At step k with $X_k = v$,

$$X_{k+1} \in \arg \min_{u \in N(v)} \text{last}[u],$$

breaking ties uniformly. Unvisited neighbors (with value 0) are preferred. Local LRV-type policies are well studied in graph exploration; they can improve practical coverage [48].

3.4. Safe Fragment Count Estimation

Because compromised relays may eavesdrop on raw key material, the source emits M fragments and estimates how many remain safe under the exposure model below. Let $\chi_{s,t}$ denote the pair-specific single-relay *exposure*: the max probability that the adversary, placed on the worst admissible intermediate relay, observes a fragment travelling $s \rightarrow t$. Because nodes do not know the topology in advance, they fix a design threshold τ and choose M from $g_\alpha^*(M, \tau)$, requiring $\chi_{s,t} \leq \tau$ for every pair (s, t) in the relevant admissible cartel class. The admissible value of τ depends greatly on the random walk variant. For a single relay v or an unordered cartel of m intermediate relays $C \subseteq V \setminus \{s, t\}$ with $|C| = m$, define the hit probability

$$p_C^{s,t} := \Pr[X \cap C \neq \emptyset].$$

Let $[A]$ denote the indicator of a condition A , equal to 1 when A is true and 0 otherwise. Empirically, for W sampled walks $X^{(1)}, \dots, X^{(W)}$, the cartel-hit probability is estimated as follows:

$$\hat{p}_C^{s,t} = \frac{1}{W} \sum_{i=1}^W [X^{(i)} \cap C \neq \emptyset].$$

Equivalently, $\hat{p}_C^{s,t}$ is the fraction of walks on which at least one relay in C is visited. For computation speed-up, the same union-hit count can be computed from marginal and intersection counts by inclusion–exclusion. For relays $u, v \in C$ let N_u be the number of walks visiting u and $N_{u,v}$ the number visiting both u and v ; for $|C| = 2$,

$$|\{i : X^{(i)} \cap C \neq \emptyset\}| = N_u + N_v - N_{u,v},$$

and for $|C| = 3$ the same idea extends to $N_a + N_b + N_c - N_{a,b} - N_{a,c} - N_{b,c} + N_{a,b,c}$. We denote the corresponding worst-case m -relay exposure over non-cut placements by

$$\chi_{s,t}^{(m)} := \max_{C \in \mathcal{D}_{s,t}^{(m)}} p_C^{s,t}, \quad \chi_{s,t} := \chi_{s,t}^{(1)}.$$

$$\mathcal{D}_{s,t}^{(m)} := \{C \subseteq V \setminus \{s, t\} : |C| = m, s \text{ and } t \text{ remain connected in } G - C\}$$

For $|C| = 1$, a relay v lies in $\mathcal{D}_{s,t}^{(1)}$ if removing v does not disconnect s from t . Assume we transmit M fragments and let G denote the number of *safe* fragments, i.e., fragments not observed by the compromised relay or relay cartel. Then, $\Pr[G \geq g]$ follows a cumulative binomial distribution (Figure 3) and can be calculated as follows:

$$\Pr[G \geq g] = \sum_{k=g}^M \binom{M}{k} (1 - \chi)^k \chi^{M-k}$$

Let us define $g_\alpha^*(M, \chi)$ as the largest g such that $\Pr[G \geq g]$ is at least, e.g., $\alpha = 99.99\%$. It is a statistical lower bound on the number of fragments that retain full entropy in transmission

$s \rightarrow t$, with failure probability at most $1 - \alpha$ under the independent-fragment exposure model and the chosen value of $\chi_{s,t}$.

$$g_{\alpha}^*(M, \chi) = \max\{g \in \{0, \dots, M\} : \Pr[G \geq g] \geq \alpha\}$$

See Table 1 for g_{α}^* values corresponding to different assumed χ values. For $M = 1024$ and $\chi = 98\%$, the table gives $g_{99.99\%}^*(1024, 98\%) = 6$. We take $\tau = 98\%$ as the default design threshold. On the synthetic 99-node graph, the worst-case HS single-relay exposure is 93.3% (Figure 4), so the recommended τ exceeds the observed single-relay exposure there; multi-node cartel coverage at this threshold is evaluated in Section 4.7. On the 43-node GÉANT topology, the same threshold also upper-bounds the worst-case HS exposure (91.3%). For a fully conservative calculation on another topology, one should use that topology's observed or assumed HS design exposure instead. The g_{α}^* calculation is deliberately conservative because it sets the fragment-count parameter M used in the information-theoretic extraction claim of Section 3.7.

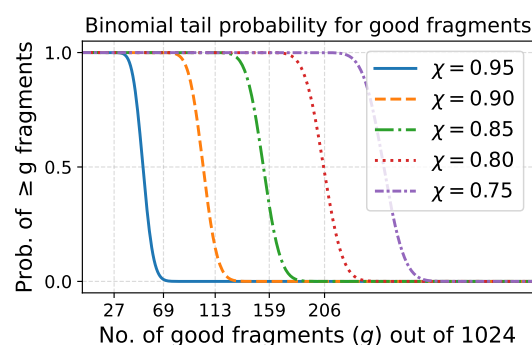


Figure 3. Probability of receiving at least g safe fragments.

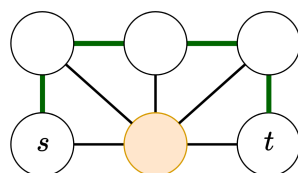


Figure 4. High χ construction example. Assume the yellow node at the bottom, between s and t , is compromised. The green line indicates the remaining safe path from s to t . Along this path, at each v , the correct neighbor must be chosen.

3.5. Computational Fallback

The exposure threshold τ determines whether the chosen output length θ has the information-theoretic extraction guarantee of Theorem 1. It does not by itself determine whether the transmission is completely disclosed. If the realized number of unseen full-entropy fragments is positive but below the lower bound needed for the configured θ , the leftover-hash guarantee for that θ no longer applies, but the adversary still has not learned the whole received fragment string. This motivates a weaker, strictly non-ITS fallback. In addition to seeded privacy amplification, an implementation can also pass the full received fragment string through a preimage-resistant 256-bit hash or KDF, such as SHA3-256 or HKDF. An adversary who has observed all but at least one full-entropy 256-bit fragment then faces a preimage/search problem over the unseen fragment bits; in the idealized model, Grover search leaves roughly 2^{128} quantum work for a 256-bit unknown [49]. This is not an information-theoretic claim, and the evaluation tables do not count it as ITS output. It only states residual computational confidentiality when at least one fragment escapes the compromised relay or relay cartel. Under the independent-fragment

exposure model, the probability that no fragment escapes is χ^M . Equivalently, over N independent transmissions, the expected number of transmissions with complete fragment disclosure is $N\chi^M$. For the default $M = 1024$, exposure $\chi = 98\%$ is therefore a useful reference point: $10^9 \cdot 0.98^{1024} \approx 1.04$. Thus, at roughly one billion transmissions, an exposure of 98% corresponds to about one expected complete-disclosure event, while lower exposure gives a rapidly smaller rate. This is the main reason for keeping M large even when the ITS-oriented extracted block θ is small: a high fragment count sharply suppresses the probability that a high-exposure relay placement observes every fragment. There is a separate defense-in-depth fallback at the application layer. As discussed in Section 2, the communicating applications can run their own post-quantum or hybrid key exchange and combine it with QKD-derived material. That endpoint mechanism does not rely on trusting the QKD relay network, whereas the fallback in this subsection is the residual confidentiality left by random flow itself when at least one fragment remains hidden.

3.6. Realized-Path Cartel Accounting

The binomial value g_α^* is a conservative ex-ante lower bound on the number of safe fragments. If verified realized paths are available, the destination can instead compute the safe-fragment count directly from those paths. How to make the destination cryptographically confident in the reported paths is outside this paper's scope and left as future work. We only record the accounting that such a mechanism would enable. For every possible compromised relay, count how many received fragments did *not* pass through that relay. The adversary is assumed to choose the relay that gives the smallest such count, so the exact worst-case single-relay count is

$$\hat{g}_{s,t}^{(1)} = \min_{v \in \mathcal{D}_{s,t}^{(1)}} \left| \{i : v \notin X^{(i)}\} \right|.$$

For an m -relay cartel, we perform the same calculation over candidate relay sets: for each cartel $C \in \mathcal{D}_{s,t}^{(m)}$, count the fragments whose paths avoid all relays in C , then take the smallest count. The corresponding value is

$$\hat{g}_{s,t}^{(m)} = \min_{C \in \mathcal{D}_{s,t}^{(m)}} \left| \{i : X^{(i)} \cap C = \emptyset\} \right|.$$

The minimum is necessary because the adversary is placed in the worst candidate relay or cartel, matching the definition of exposure as a maximum hit probability. Define the realized-path safe-fragment yield as $\rho_{s,t}^{(m)} := \hat{g}_{s,t}^{(m)} / M$. Under the independent-fragment exposure model, for the universal target $m = \kappa(s, t) - 1$,

$$\rho_{s,t}^{(m)} \xrightarrow{M \rightarrow \infty} 1 - \chi_{s,t}^{(m)}.$$

After the two-fragment privacy-amplification penalty below, the extracted-key yield $\rho_{s,t}^{\text{ext}} := \theta_{s,t} / M$ has the same asymptotic limit.

3.7. Privacy Amplification

Let $\theta_{s,t}$ denote the number of final 256-bit keys extracted after privacy amplification for transmission $s \rightarrow t$. Then, g_α^* should be understood as a lower bound on the number of safe fragments, not directly on $\theta_{s,t}$. Since verified realized paths are left to future work, the default privacy-amplification parameter is derived from the configured design threshold τ :

$$g_\tau := g_\alpha^*(M, \tau), \quad \theta_\tau := \max(0, g_\tau - 2).$$

Here, g_τ is a high-confidence lower bound on the number of safe fragments under the independent-fragment exposure model and the condition that every admissible cartel has per-fragment exposure at most τ . If verified realized paths become available, the exact count $\hat{g}$ from Section 3.6 can replace g_τ . At the destination, the recovered fragment material is concatenated into Z and processed by a seeded universal-2 extractor, $K = \text{Ext}_S(Z)$, where the seed S is public but authenticated end to end. For an ITS claim, this step should use a privacy-amplification construction covered by the Leftover Hash Lemma, for example a universal-hash implementation such as Toeplitz hashing [50]. It should not be replaced by an ordinary cryptographic hash function when the claim being made is information-theoretic security. Cryptographic hashes or KDFs only support the computational fallback discussed in Section 3.5. Given g_τ , we can state the following theorem.

Theorem 1. *Assume that, for transmission $s \rightarrow t$, each admissible relay cartel observes any one fragment with probability at most τ , and fragment exposure events are independent across fragments. Let $g_\tau = g_\alpha^*(M, \tau)$. Then seeded universal-2 privacy amplification with $\varepsilon_{\text{PA}} = 2^{-256}$ extracts $\theta_\tau = \max(0, g_\tau - 2)$ full 256-bit keys that are ε_{PA} -close to uniform except with probability at most $1 - \alpha$ under the stated exposure model. No extraction occurs when $g_\tau \leq 2$.*

Proof. Fix the public transcript (extractor seed and any public protocol messages) and an admissible adversary $\mathcal{A}$. By the definition of $g_\tau = g_\alpha^*(M, \tau)$, the event that $\mathcal{A}$ misses at least g_τ of the M fragments has probability at least α under the independent-fragment exposure model. Condition on this event. Each fragment $f_i \in \{0, 1\}^{256}$ is sampled at the source uniformly and independently of all other fragments and of the public transcript. Therefore, conditioned on the public transcript and $\mathcal{A}$'s view of the at most $M - g_\tau$ fragments it observes on this event, the remaining g_τ unseen fragments are jointly uniform on $\{0, 1\}^{256g_\tau}$. The conditional min-entropy of the M -fragment string given $\mathcal{A}$'s view is therefore at least $256g_\tau$ bits. By the Leftover Hash Lemma with a universal-2 hash family and an independent public seed [51–53], privacy amplification extracts at most

$$\ell \leq 256g_\tau - 2\log_2(1/\varepsilon_{\text{PA}})$$

bits whose distribution is ε_{PA} -close to uniform conditioned on $\mathcal{A}$'s view. Flooring to whole 256-bit keys gives

$$\theta_\tau = \max\left(0, \left\lfloor \frac{256g_\tau - 2\log_2(1/\varepsilon_{\text{PA}})}{256} \right\rfloor\right).$$

For $\varepsilon_{\text{PA}} = 2^{-256}$, this reduces to $\theta_\tau = \max(0, g_\tau - 2)$. The seed can be public because LHL is a strong extractor; in our setting, it is transmitted over the authenticated classical channel so that source and destination derive the same keys. The event $G \geq g_\tau$ fails with probability at most $1 - \alpha$, and any residual error in the assumed exposure bound τ contributes an additional term ε_{est} . By a union bound, the total distinguishing advantage between the extracted key and an ideal uniform key, conditioned on the public transcript and adversarial view, satisfies

$$\varepsilon_{\text{sec}} \leq (1 - \alpha) + \varepsilon_{\text{PA}} + \varepsilon_{\text{est}}.$$

□

3.8. Exposure Reduced RW Variants

The simulation results in Section 4.2 show that the fully local variants R, NB, and LRV have high $\max_{s,t} \chi_{s,t}$ values. On each evaluated graph, there exists a pair s, t such that $\chi_{s,t} \geq 96.4\%$ for each of these walk variants. See Figure 4 for intuition about why high χ

can arise. There is often a longer bypass path connecting s and t , but the walk is likely to be diverted back into the more “congested” direct connection. To address this, we can temporarily color nodes one by one, or otherwise assign them lower priority throughout the walk. Different walks within the same transmission should penalize different nodes. The following constructions reduce $\max \chi$ from 95.3% (LRV) to 91.3% (HS). The expected number of fragments not observed by single-relay cartel is proportional to $1 - \chi$. It therefore increases by 85.1% because $1 - 95.3 = 4.7\%$, and $1 - 91.3 = 8.7\%$.

One-by-one node-coloring (NC). If s and t are biconnected, it is possible to exclude nodes one-by-one from the graph, by coloring (marking) them. That still keeps connectivity between s and t but forces the random walk to avoid the excluded vertex, thus guaranteeing that it will not know at least one fragment. This approach, however, requires the node-identifier universe W from which excluded relays are selected. Thus, NC should be described as *partially topology-oblivious*. Nevertheless, the NC variant does not require adjacency knowledge, or path computation. If $M > |W|$, we can color vertices in circular order (thus, some nodes will miss several fragments). Otherwise, when choosing the next hop, we apply the LRV walk strategy, taking into a consideration the colored (excluded) node.

Theorem 2. *In the presence of a single malicious node, the NC random walk variant guarantees at least one key that is uniform under this idealized secret-sharing model.*

Proof. This is the standard additive (XOR) secret-sharing argument [54]. For $K \in \{0, 1\}^{256}$, sample $f_1, \dots, f_{M-1}$ uniformly and set $f_M = K \oplus \bigoplus_{i=1}^{M-1} f_i$, so $K = \bigoplus_{i=1}^M f_i$. If the compromised relay misses at least one fragment, then the XOR of the fragments it misses is still uniform over $\{0, 1\}^{256}$ and independent of the fragments it sees. XORing this unknown uniform value with the relay’s fixed view leaves K uniform from that relay’s perspective. $\square$

Although additive (XOR) secret sharing is easy to implement and reason about, we instead quantify χ and apply entropy extraction (as in Theorem 1) for higher yield ρ .

Highest-score neighbor (HS). HS is a seed-based local diversification heuristic. For each fragment token i , the source samples a fresh seed $\zeta_i \xleftarrow{\$} \{0, 1\}^\lambda$. The seed gives each vertex u a deterministic per-token score

$$\text{score}_i(u) = h(u, \zeta_i),$$

where h is a deterministic mixing function of the vertex identifier u and seed ζ_i . Different seeds induce different local rankings without requiring topology knowledge. Thus, a high-exposure relay can receive a low score for some fragments. Let $U_i(v) \subseteq N(v)$ be the neighbors not yet visited by token i . If $U_i(v)$ is nonempty, HS chooses a highest-scoring unvisited neighbor, breaking ties uniformly:

$$X_{k+1}^{(i)} \in \arg \max_{u \in U_i(v)} \text{score}_i(u)$$

If all neighbors have already been visited, HS falls back to the NB rule. Operationally, a relay computes these scores only for its current neighbors and for the current token seed ζ_i . No relay stores a persistent score table; the token carries only its seed, previous hop, and visited-node set. Figure 5 illustrates how different per-fragment seeds can rank the same relay differently and thereby diversify exposure.



Figure 5. Two HS score assignments on the same topology. Node colors indicate the per-seed score ordering, and the highlighted lines indicate the selected next-hop choices. In the first run, the central relay has the lowest score and is avoided; in the second, it has a higher score than one of the nodes and is selected.

3.9. Efficiency and Throughput

For a fixed source–destination pair (s, t) , let $H_{s,t}$ denote the random hop count of one fragment walk, and let $h_{s,t} := \mathbb{E}[H_{s,t}]$ be the expected hop count. As previously assumed, the fragment size is 256 bits and matches link key size. One delivered fragment therefore consumes, in expectation, $h_{s,t}$ link keys. A transmission consumes $M \cdot h_{s,t}$ link keys. Suppose a transmission uses M fragments. Under the independent-fragment exposure model, the expected fraction of fragments that avoid a worst-case admissible cartel is the complement of the corresponding cartel exposure. Let $m := \kappa(s, t) - 1$. For a fixed design threshold τ on admissible cartel exposure, define the pair-specific expected yield and the conservative design yield side by side:

$$\rho_{s,t} := 1 - \chi_{s,t}^{(m)}, \quad \rho^\tau := \frac{g_\alpha^*(M, \tau)}{M}.$$

Section 3.6 motivates this choice: $\hat{g}_{s,t}^{(m)} / M$ converges to $1 - \chi_{s,t}^{(m)}$, so $\rho_{s,t}$ is the expected safe-fragment count per emitted fragment, not the conservative binomial lower bound $g_\alpha^*(M, \chi) / M$ used for ITS parameter selection. Thus, ρ^τ is a high-confidence lower bound on the safe-fragment yield under the design threshold, while $\rho_{s,t}$ is the pair-specific expected yield. We use ρ throughout for safe-fragment yield and reserve θ for the extracted-key count, whose default value is $\theta_\tau = \max(0, g_\tau - 2)$ from Theorem 1. Because QKD-derived link keys are scarce, we define the *model-based extracted efficiency* as a heuristic proxy for extracted output bits per consumed QKD-derived link-key bit:

$$\eta_{s,t} := \rho_{s,t} h_{s,t}^{-1}, \quad \eta_{s,t}^\tau := \rho^\tau h_{s,t}^{-1}.$$

This is not an exact extracted-key count: ρ tracks safe-fragment yield before the additive LHL penalty, while exact default counts use $\theta_\tau = \max(0, g_\tau - 2)$. For the reported $M \in \{256, 1024\}$, this penalty is at most 0.8% of M and is included only in quantitative claims about θ . When reporting a topology-wide metric over a set $\mathcal{P}$ of ordered (s, t) pairs, we use

$$\bar{\eta} := \frac{1}{|\mathcal{P}|} \sum_{(s,t) \in \mathcal{P}} \eta_{s,t}, \quad \bar{\eta}^\tau := \frac{1}{|\mathcal{P}|} \sum_{(s,t) \in \mathcal{P}} \eta_{s,t}^\tau.$$

Next, define the raw fragment throughput $T_{s,t}$ as the long-run average rate, in bits/s, at which fragment bits arrive at the destination t . We assume that only one ordered pair (s, t) is active at a time, that QKD links continuously generate and buffer keys (starting with empty buffers). The corresponding model-based extracted throughput is given by

$$R_{s,t} := T_{s,t} \rho_{s,t}, \quad \text{equivalently } R_{s,t} = T_{s,t} \eta_{s,t} h_{s,t}.$$

Thus, $\eta_{s,t}$ captures model-based extracted efficiency, while $R_{s,t}$ captures the corresponding extracted-throughput proxy under the simplified entropy model.

3.10. Loop Erasure

We separate route discovery from payload transmission. In the first phase, the source emits a lightweight *scouting token* over the authenticated classical channel only. The scouting token carries no secret fragment material and consumes no QKD-derived link-key bits. It is forwarded according to the same stochastic forwarding rule as the corresponding payload walk and records the visited history

$$W = (v_0 = s, v_1, \dots, v_\tau = t).$$

When we later report exposure for loop-erased variants, we still measure it on this original scouting walk as a conservative upper bound on literal payload exposure. During the walk a time-bounded link key reservation is applied to traversed links. Ultimately, if some relay along this sampled walk cannot admit the request, the session may be rejected before any key material is transmitted; in a REST realization, this can be surfaced as an implementation-level failure such as HTTP 503. After the scouting token reaches t , the recorded walk W is converted into a realized payload route by *chronological loop erasure* [55,56]. Concretely, one scans W from left to right and, whenever a vertex reappears, deletes the closed subwalk between the earlier occurrence and the repeat, while retaining a single copy of the repeated vertex. Let $\text{LE}(W)$ denote the resulting simple s – t path. In the second phase, the actual fragment material is forwarded hop by hop along $\text{LE}(W)$, with each hop protected by fresh QKD-derived link-key bits used as a one-time pad. Strictly speaking, for the NB, LRV, NC, and HS variants, $\text{LE}(W)$ is not the *classical* loop-erased random walk distribution from probability theory; here loop erasure is used only as a path-simplification operator applied to the sampled scouting history. This two-phase construction provides two operational benefits. First, removing loops shortens the realized payload route, thereby reducing expected hop count, QKD key consumption, and queueing pressure. Second, because the realized payload route is simple, self-induced cyclic waiting caused by one fragment revisiting the same relays is eliminated. Rejection, if necessary, occurs before secret material enters the relay path. The interpretation of exposure under loop erasure requires care. A compromised relay that appears only on an erased segment of the scouting walk does not directly observe fragment plaintext, because the scouting phase is classical only. However, such a relay may still influence the realized route by biasing the scouting trajectory or by affecting admission decisions. Therefore, when reporting $\chi_{s,t}$ for loop-erased variants, we count scouting hits as a *conservative upper bound* on literal payload exposure. This preserves comparability with the no-loop-erasure case while avoiding an optimistic estimate of adversarial influence.

3.11. Network Dynamism

The topology-oblivious forwarding model does not remove the need for controlled topology changes. Relevant changes include adding a relay or QKD link, removing a relay and its incident links, and removing a QKD link. We assume supervised changes: the administrator of a relay manually registers or unregisters its incident QKD links and neighboring relays. Because the operation is local, the administrator does not need a full network view. Adding QKD links cannot decrease the minimum vertex cut between any fixed source–target pair [46]. It can nevertheless change exposure by making some relays more likely to appear on sampled walks, so exposure should be re-evaluated after topology expansion. Adding a relay can also reduce throughput until the new incident QKD links have enough key material, because pairwise link keys must be established with that relay. If the new relay is later compromised, fragments traversing it no longer contribute full entropy, and the privacy-amplification output length must be reduced according to the updated exposure model. We argue against automatic removal of QKD nodes or links from

the forwarding graph. Silent removal can change the security model seen by forwarding relays without updating the configured exposure threshold. For brittle QKD-link failures or exhausted link-key buffers, the safer behavior is fail-stop rejection. If a relay cannot forward using an available QKD link while satisfying the configured threat model and key-availability assumptions, the transmission should fail and the operator should be notified. This reduces availability, but it prevents the implementation from silently routing under an exposure model that no longer matches the configured parameters. Classical-link dynamism is treated separately. The classical transport stack retransmits lost packets and eventually fails after a timeout. Such failures are reported to the operator, who can then reconfigure the QKD-network view if needed.

4. Simulation Study

We focus on three questions: (i) *exposure*—how large the worst-case relay hit probability $\chi_{s,t}$ can be for some source–target pair, and what variant-specific worst-case or illustrative design exposure this suggests when choosing the target extracted block size θ and, where applicable, the fragment count M ; (ii) *efficiency*—what extracted-yield estimators and model-based output per consumed QKD-derived link-key bit the evaluated variants achieve under the paper’s simplified entropy model; and (iii) *scalability*—how these quantities and the routing cost evolve with network size. We address these questions on topologies resembling existing quantum network deployments.

4.1. Simulated Topologies

Trusted-node QKD deployments are currently small (typically tens of nodes) due to the cost of dark fiber and QKD equipment. For context, the largest openly reported QKD deployment is the China Quantum Communication Network (CN-QCN), spanning 145 nodes [57]; however, CN-QCN uses centralized network management and exhibits a hierarchical topology with many articulation points. For our simulation study, we selected two deployed-analog topologies with 14 and 43 nodes (NSFNET and GÉANT). To study scaling trends, we additionally use a synthetic 99-node graph described in Section 4.8. Both topologies are visualized in Gephi 0.10 using a geographic layout based on approximate site latitude/longitude, and node colors denote Gephi modularity communities with no physical or administrative meaning. We deliberately focus on graphs with a nontrivial biconnected structure, and, whenever we report ITS-oriented exposure, yield, efficiency, or throughput quantities, we restrict attention to biconnected ordered pairs (s, t) . The NSFNET T1 backbone from 1991 with 14 nodes [58]. Although NSFNET is a classical network, it serves as a medium-sized, historically relevant analog: in early wide-area optical transport, capacity was scarce and expensive, motivating sparse topologies and careful end-to-end provisioning. Similar constraints reappear in QKD networks. Figure 6 shows the reconstructed graph used in the simulations.

The GÉANT GN4 Phase 3 backbone (GN4-3N) is a large, well-connected topology. Our variant, after pruning links longer than 1000 km to improve topology perceptibility in the graph diagram, has 43 nodes and 59 edges. GÉANT is also engaged in Europe’s “ultra-secure” communications direction (including EuroQCI-oriented efforts that consider QKD overlays), making this backbone a plausible substrate for a future QKD overlay [59,60]. Figure 7 shows the adapted GÉANT graph.

Table 2 summarizes basic structural metrics of the two deployed-analog topologies.

In practice, the GÉANT and NSFNET topologies would require many additional relay sites because their longest links exceed the ≈ 150 –300 km range typical of current commercial devices. We do not introduce such relays here; these two graphs are treated

as structural proxies. Notice also that, as pointed out in the introduction, dense QKD networks are not expected in the near future.

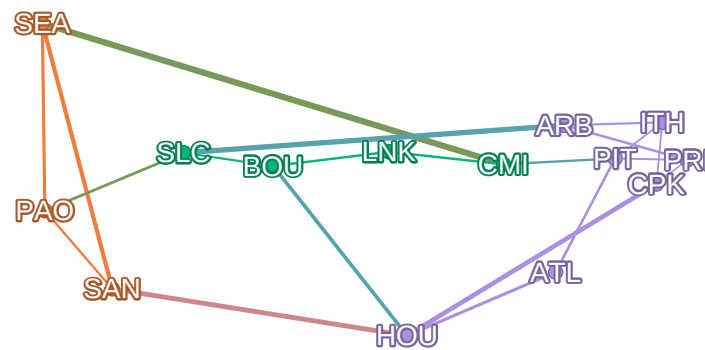


Figure 6. Reconstructed NSFNET T1 topology. Node colors denote Gephi modularity communities and have no physical or administrative meaning.

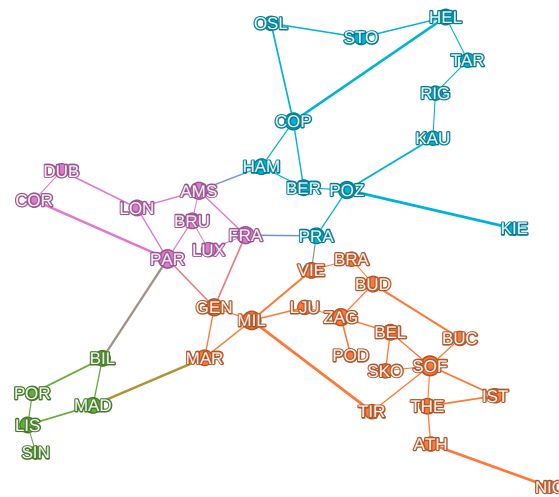


Figure 7. Adapted version of GÉANT GN4-3N. Node colors denote Gephi modularity communities and have no physical or administrative meaning.

Table 2. Topology overview. ASP is the average shortest-path length. The final columns report the percentage of ordered endpoint pairs with vertex connectivity $\kappa(s, t)$ at least 2 or 3.

Graph	Nodes	Edges	Diam.	Avg. Deg.	ASP	$\kappa \geq 2$	$\kappa \geq 3$
NSFNET	14	21	3	3.00	2.143	100%	72.5%
GÉANT	43	59	12	2.74	4.682	70.1%	5.0%

4.2. Single-Node Exposure

To evaluate the security-relevant exposure, we estimate the exposures defined in Section 3.4 by Monte Carlo sampling. Throughout this section, hat notation marks empirical quantities computed from sampled walks; for example, $\hat{p}_v^{s,t}$ and $\hat{\chi}_{s,t}$ are the sampled counterparts of $p_v^{s,t}$ and $\chi_{s,t}$. For single-relay exposure, we apply the cartel-hit estimator there with $|C| = 1$, i.e., $\hat{p}_v^{s,t} := \hat{p}_{\{v\}}^{s,t}$ and $\hat{\chi}_{s,t} = \max_{v \in \mathcal{D}_{s,t}^{(1)}} \hat{p}_v^{s,t}$. We sample $W = 10,000$ loop-erased walks per ordered source–target pair and variant. Each walk uses the run index $r \in \{0, \dots, W - 1\}$ as its deterministic PRNG seed in the C++ exposure simulator, compiled with g++ 16.1.1 using default compile settings; hop-count experiments use the same convention with $W = 1000$. The synthetic scaling graph is generated with Python 3.12 seed 2026. We report point estimates only: at $W = 10,000$ the Monte Carlo standard error on a 90% hit rate is about 0.3 percentage points, so small differences be-

tween close variants should not be overinterpreted. Graph files, reconstruction scripts, and analysis code are available in the accompanying repository. In each graph, we then identify $\max_{s,t} \hat{\chi}$ for each RW variant. For the single-relay admissible class, this worst-case value is the relevant exposure to compare with the design threshold τ when setting M and the extracted block size θ_τ through $g_\alpha^*(M, \tau)$. For loop-erased variants, $X^{(i)}$ still denotes the original sampled/scouting walk rather than only the realized payload route, so the reported $\hat{\chi}$ remains a conservative upper bound on literal payload exposure. Table 3 reports the detailed GÉANT breakdown, including the maximizing (s, t, v) triple and the average $\hat{\chi}$ over all ordered pairs, while Table 4 summarizes per-graph max and median $\hat{\chi}$ across variants. On GÉANT, NC exceeded the walk step limit on 530 ordered pairs; those pairs are omitted from the NC aggregates in Tables 3 and 4, which is one reason NC is treated as a diagnostic variant rather than an implementer-facing recommendation.

Table 3. RW variant $\max_{s,t} \hat{\chi}$ values (in %) on GÉANT and corresponding intermediate vertex v under the single-relay metric.

Variant	Max $\hat{\chi}$	s	t	v	Avg $\hat{\chi}$	Median $\hat{\chi}$
R	98.3	COR	BIL	PAR	71.0	78.1
NB	95.7	BIL	COR	PAR	66.9	72.7
LRV	95.3	POR	COR	PAR	67.0	73.7
NC	93.6	POR	COR	PAR	65.8	72.8
HS	91.3	MAD	COR	PAR	65.0	71.7

Table 4. Exposure overview by graph and RW variant under the single-relay metric. For each graph and variant, we report the worst-case pair exposure $\max_{s,t} \hat{\chi}$ and the median $\hat{\chi}$ over all ordered source–target pairs.

Graph	Max Exposure $\hat{\chi}$ [%]					Median Exposure $\hat{\chi}$ [%]				
	R	NB	LRV	NC	HS	R	NB	LRV	NC	HS
NSFNET	81.2	78.1	77.4	72.8	71.0	52.6	50.9	52.8	50.8	51.6
GÉANT	98.3	95.7	95.3	93.6	91.3	78.1	72.7	73.7	72.8	71.7

As reported in Tables 3 and 4, HS attains the lowest worst-case exposure on NSFNET, and also the lowest worst-case exposure on GÉANT under the non-cut single-relay metric. The reported maximum is a max-of-max quantity: for each ordered source–target pair, the relay position is chosen adversarially within $\mathcal{D}_{s,t}^{(1)}$, and the largest such value is then reported over all pairs. On GÉANT, the worst eligible single relay for HS is PAR on the MAD–COR pair, giving $\hat{\chi} = 91.3\%$. The median exposure is lower (71.7% for HS), meaning that for a typical ordered pair, even the best-positioned non-cut compromised relay misses roughly a quarter of the fragments. Among the evaluated variants, HS gives the best worst-case exposure. NC remains close on GÉANT, but it assumes the globally known node set, whereas HS preserves the stronger locality of using only neighbor information plus token state.

4.3. Multi-Node Exposure

We also use the same empirical exposure calculation to evaluate two- and three-node compromised relay cartels C . Call an ordered pair (s, t) *eligible* for cartel C if $s, t \notin C$ and s and t remain connected in the subgraph $G[V \setminus C]$ obtained by removing all cartel nodes—equivalently, C is not an s – t vertex cut, so at least one C -disjoint s – t path exists in $G[V \setminus C]$ and a topology-aware multipath baseline (Section 4.7) could in principle route safely between them. Formally,

$$\mathcal{E}(C) = \{(s, t) : s, t \notin C, s \text{ and } t \text{ are connected in } G[V \setminus C]\}.$$

A cartel is itself called *eligible* if $|\mathcal{E}(C)| > 0$, i.e., it does not disconnect every source–target pair. Call an eligible pair (s, t) *affected* at design threshold τ if its empirical cartel exposure exceeds that threshold:

$$\mathcal{A}(C, \tau) = \{(s, t) \in \mathcal{E}(C) : p_C^{s,t} > \tau\}.$$

An affected pair is one for which the chosen threshold τ would overestimate the number of full-entropy output keys for that cartel. This does not necessarily mean that every fragment is exposed: as discussed in Section 3.5, if a cartel’s per-fragment exposure is $\chi_C < 1$, then the independent-fragment model gives probability χ_C^M that the cartel observes all fragments. For example, $\chi_C = 99\%$ and $M = 1024$ gives $0.99^{1024} \approx 0.0034\%$. Hence, even when a cartel invalidates the configured ITS parameter choice, the strictly non-ITS computational fallback discussed in Section 3.5 (preimage-resistant 256-bit hash or KDF over the received fragment string) still applies whenever at least one full-entropy fragment is not observed. This is not the information-theoretic claim; it is only a computational fallback when at least one full-entropy fragment remains hidden. In the idealized model, Grover search leaves roughly 2^{128} quantum work for a 256-bit unknown [49]. Table 5 reports this GÉANT HS analysis for two- and three-node cartels. For each candidate design threshold τ and cartel size, it gives the median, average, and maximum of $|\mathcal{A}(C, \tau)| / |\mathcal{E}(C)|$ over all eligible unordered cartels C .

Table 5. Two- and three-relay compromise analysis on GÉANT using HS. For each design threshold τ , the table reports the fraction of eligible pairs that are *affected* ($|\mathcal{A}(C, \tau)| / |\mathcal{E}(C)|$), over all cartels C .

τ	Two-Node Cartel			Three-Node Cartels		
	Median	Average	Maximum	Median	Average	Maximum
75%	0.61%	5.56%	60.98%	7.31%	13.24%	68.42%
80%	0.00%	3.17%	54.94%	2.88%	8.76%	65.65%
85%	0.00%	1.44%	41.83%	0.47%	4.92%	60.93%
90%	0.00%	0.41%	31.65%	0.00%	2.00%	53.17%
95%	0.00%	0.03%	5.79%	0.00%	0.34%	33.47%
97%	0.00%	0.00%	0.79%	0.00%	0.07%	24.62%
98%	0.00%	0.00%	0.00%	0.00%	0.02%	19.49%
99%	0.00%	0.00%	0.00%	0.00%	0.00%	3.33%

These summary statistics should be interpreted descriptively: most two- or three-node cartels are not chosen to target a specific source–target pair and are therefore placed rather arbitrarily with respect to any given (s, t) . At the default $\tau = 98\%$, the average affected-pair fraction is 0.00% for two-node cartels and 0.02% for three-node cartels. The maximum column is nevertheless useful as a warning about favorable placements for a relay cartel: the worst eligible two-relay cartel affects no eligible ordered pairs at $\tau = 98\%$, while the worst eligible three-relay cartel affects 19.49%. Thus, the single-relay exposure results should not be read as a multi-relay guarantee, even though a computational fallback remains when at least one fragment escapes the cartel.

4.4. Expected Hop Count

To put these numbers into context, they remain well above the corresponding average shortest-path lengths on NSFNET and GÉANT. The longest mean path after loop-erasure is attained by HS. This aligns with its original path-diversification motivation.

4.5. Efficiency

From Section 3.9, pairwise efficiency factors as $\eta_{s,t} = \rho_{s,t} h_{s,t}^{-1}$ with $\rho_{s,t} = 1 - \chi_{s,t}^{(\kappa(s,t)-1)}$. We further distinguish *pair-specific* efficiency from *design-threshold* efficiency $\eta_{s,t}^\tau = \rho^\tau h_{s,t}^{-1}$ with $\rho^\tau = g_\alpha^*(M, \tau) / M$. For the pair-specific case, we write

$$\eta_{s,t}^{\hat{\chi}} := (1 - \hat{\chi}_{s,t}^{(m)}) h_{s,t}^{-1}, \quad m := \kappa(s, t) - 1,$$

where $\hat{\chi}_{s,t}^{(m)}$ is the measured $(\kappa(s, t) - 1)$ -relay cartel exposure for (s, t) . For $\bar{\eta}^{\tau}$, we set τ to the graph- and variant-specific worst-case single-relay exposure $\max_{s,t} \hat{\chi}$ from Section 4.2 (Table 4). On NSFNET, these thresholds range from 71.0% (HS) to 78.1% (NB); on GÉANT, from 91.3% (HS) to 95.7% (NB). In the simulation results below, we report topology-wide averages $\bar{\eta}^{\tau}$ and $\bar{\eta}^{\hat{\chi}}$ over all biconnected ordered pairs. Since exposure is still computed from the original walk, loop-erasure affects only the denominator through the expected hop count. In Tables 6 and 7, we omit the simple random walk R and focus on NB, LRV, NC, and HS, since R is already dominated on the evaluated graphs in both exposure and expected hop count. Table 6 shows that the design-threshold efficiency substantially understates the pair-specific efficiency on all evaluated topologies: $\bar{\eta}^{\hat{\chi}}$ is consistently far above $\bar{\eta}^{\tau}$. The reason is that $\bar{\eta}^{\tau}$ fixes one worst-case exposure per graph and variant, whereas $\bar{\eta}^{\hat{\chi}}$ uses the exposure measured for each source–target pair. Within each metric row, Table 6 also compares original walks with loop-erased routes. The right block uses the loop-erased routes, and the left block uses the original sampled walks. The resulting increase in efficiency is noticeably smaller than the corresponding reduction in mean hop count from Table 8. This is not a contradiction. Although loop-erasure reduces average hop count dramatically, efficiency scales with the reciprocal hop count h^{-1} , not with h itself. Thus loop-erasure primarily removes a heavy tail of excessively long walks, rather than uniformly increasing h^{-1} for all source–target pairs.

The distribution of pairwise efficiencies is strongly right-skewed, especially on GÉANT. For instance, for NB without loop-erasure on GÉANT, the mean $\bar{\eta}^{\hat{\chi}}$ is 9.7% whereas the median is only 0.72%; with loop-erasure, the corresponding values are 12.1% and 3.1%. Thus, the mean efficiency should be read as a topology-wide average, not as a typical per-pair value.

Table 6. Topology-wide mean RW efficiency by graph and variant, in percent. Each row reports $\bar{\eta}^{\tau}$ (fixed τ from Table 4 and $\rho^{\tau} = g_{\alpha}^*(M, \tau) / M$) or $\bar{\eta}^{\hat{\chi}}$ (pairwise measured $\hat{\chi}_{s,t}^{(m)}$). The left block uses hop counts from the original walk, while the right block uses hop counts after loop-erasure. In both blocks, exposure is measured on the original walk and $M = 1024$.

Graph	Mean	Loop-Erasure off [%]				Loop-Erasure on [%]			
		NB	LRV	NC	HS	NB	LRV	NC	HS
NSFNET	$\bar{\eta}^{\tau}$	6.49	7.07	8.69	9.47	7.25	7.23	9.12	9.63
	$\bar{\eta}^{\hat{\chi}}$	26.5	26.8	26.9	27.1	27.6	27.1	27.4	27.2
GÉANT	$\bar{\eta}^{\tau}$	0.27	0.36	0.54	0.79	0.48	0.49	0.76	1.09
	$\bar{\eta}^{\hat{\chi}}$	9.7	10.2	10.2	10.2	12.1	11.6	11.8	11.7

Table 7. Mean throughput by graph and RW variant. The top block reports raw fragment throughput $T_{s,t}$ in kbit/s. The bottom block reports mean model-based extracted throughput $R_{s,t} = T_{s,t} \rho_{s,t}$ in kbit/s under the simplified entropy model.

Metric	Graph	NB	LRV	NC	HS
$T_{s,t}$	NSFNET	2.03	2.02	2.05	2.03
	GÉANT	1.70	1.70	1.72	1.77
$R_{s,t}$	NSFNET	0.64	0.62	0.64	0.64
	GÉANT	0.50	0.50	0.51	0.56

Table 8. Estimated expected hop count by graph and RW variant. Let $H_{s,t}$ be the hop count in the sampled walk from s to t . We sample $W = 1000$ walks per biconnected ordered source–target pair and variant and report the average of mean, median, and maximum of $H_{s,t}$ over all biconnected (s, t) pairs, distinguishing whether loop-erasure is toggled on.

Graph	Metric	Without Loop-Erasure					With Loop-Erasure				
		R	NB	LRV	NC	HS	R	NB	LRV	NC	HS
NSFNET	Mean	6.4	4.5	4.0	4.1	3.9	3.1	3.4	3.7	3.6	3.7
	Median	5	4	3	4	3	3	3	3	3	3
	Max	40	23	9	12	10	7	8	8	8	8
GÉANT	Mean	64.4	32.2	18.4	18.9	19.8	6.3	7.2	8.3	8.1	8.7
	Median	42	22	15	16	15	6	6	7	7	8
	Max	512	243	70	83	130	18	21	23	23	24

4.6. Throughput

To evaluate *throughput* for a fixed (s, t) pair, we implement the random-walk key-relaying scheme as a discrete-event simulation driven by a min-heap of timestamped events. To remain consistent with the two-phase protocol model, route discovery is treated as a preceding classical control phase that determines the realized payload route (with loop-erasure enabled, this route is $LE(W)$). The throughput metric $T_{s,t}$ counts only delivered payload fragment bits; scouting/control traffic is omitted from $T_{s,t}$ because it uses the classical channel and consumes no QKD-derived OTP key material. Each hop of a packet from node u to node v is split into: (i) *OTP key reservation* on link (u, v) , which may incur waiting time computed from the link secret-key rate (SKR) and current key balance; (ii) a fixed-latency classical channel; and (iii) *receiver admission* into a finite relay buffer with FIFO backpressure. The simulator repeatedly pops the earliest event, advances the simulation clock, updates link/node state, and schedules successor events. Each undirected link $e \in E$ generates QKD key material at a constant secret-key rate $g_e = 1$ kbit/s for all links. Key material is consumed in 256 bit units. Each hop spends 256 bit for one-time-pad (OTP) encryption on the traversed QKD link. We use a default inter-node classical latency of 5 ms. Simulations start with empty link buffers. We do not discard a warm-up period when calculating throughput; given a fixed simulation duration of 1000 s, its impact is unlikely to be large. We ignore classical network throughput (goodput) and protocol overhead, as QKD is the dominant bottleneck (e.g., 1 kbit/s vs. typical 1 Gbit/s classical links). Table 7 summarizes both the mean raw fragment throughput and the mean model-based extracted throughput across graphs and RW variants.

Raw fragment throughput varies only modestly across variants, but extracted throughput differs more because it inherits the pair-specific yield $\rho_{s,t}$. On NSFNET, mean raw throughput is about 2.0 kbit/s for all evaluated variants, whereas mean model-based extracted throughput ranges from roughly 0.62 to 0.64 kbit/s. On GÉANT, mean raw throughput ranges from about 1.70 to 1.77 kbit/s, while mean model-based extracted throughput rises from about 0.50 kbit/s for NB to about 0.56 kbit/s for HS. If the default design threshold $\tau = 98\%$ is used instead of pair-specific $\rho_{s,t}$, then $\rho^\tau = 6/1024$ and the corresponding mean extracted-throughput proxy is only about 0.01 kbit/s. This mirrors the conservative-design efficiency collapse discussed after Table 9.

Table 9. Pair-mean efficiency comparison. RF column reproduces $\bar{\eta}^\chi$ from Table 6 (HS, with loop-erasure, $M = 1024$). MP column reports $\bar{\eta}^{MP,2} = 1/(L_1 + L_2)$ averaged over biconnected ordered pairs using Suurballe shortest disjoint-pair paths.

Graph	$\bar{\eta}_{HS}^\chi$ [%]	$\bar{\eta}^{MP,2}$ [%]	RF/MP
NSFNET	27.3	18.1	1.5
GÉANT	11.8	11.1	1.1

4.7. Comparison with Baseline

To put random flow on common footing with a topology-aware deterministic alternative, we introduce a node-disjoint multipath (MP) baseline. For each biconnected ordered pair (s, t) , MP selects a set of k node-disjoint s – t paths $\{P_1^{s,t}, \dots, P_k^{s,t}\}$ of minimum total length via Suurballe’s algorithm [27,28], which relies on repeated shortest-path computations (Dijkstra-class routines) over the known topology. Random-flow forwarding deliberately avoids such end-to-end shortest-path or disjoint-path computation at the source. The source XOR-splits each 256-bit final key K into k uniformly random 256-bit shares with $K = S_1 \oplus \dots \oplus S_k$, transmits share S_i along $P_i^{s,t}$ with each hop protected by a fresh QKD-derived 256-bit OTP key, and the destination reconstructs K as $\bigoplus_{i=1}^k S_i$. The secrecy argument is the same as in Theorem 1, restricted to k shares across disjoint paths instead of M fragments along random walks. Under the same single-compromise threat model used for random flow, any compromised relay lies on at most one of the k disjoint paths and therefore observes at most one of the k shares. Since $k - 1$ shares are needed to recover K in the information-theoretic sense, the worst-case MP fragment-exposure for any biconnected pair is exactly $1/k$. Because MP is topology-aware, the source picks the largest feasible $k = \kappa(s, t)$ per pair: with $k = \kappa(s, t) \geq 2$ on every biconnected pair, the worst-case fragment exposure under a single compromise is $1/\kappa(s, t) \leq 50\%$ uniformly. Let $L_i^{s,t} = |P_i^{s,t}|$ denote the hop count of the i -th chosen disjoint path. One MP execution consumes $\sum_{i=1}^k L_i^{s,t}$ QKD link-key blocks of 256 bits and produces one 256-bit final key. The model-based efficiency, in the same units as $\eta_{s,t}$ (extracted output bits per consumed QKD-derived link-key bit), is therefore

$$\eta_{s,t}^{\text{MP},k} := \frac{1}{\sum_{i=1}^k L_i^{s,t}}.$$

This is the natural pairwise analogue of $\eta_{s,t} = \rho_{s,t} h_{s,t}^{-1}$ for the deterministic baseline: $\rho^{\text{MP},k} \equiv 1/k$ (one final key per k transmitted shares) and the aggregate hop count is $\sum_i L_i^{s,t} = k \bar{L}^{s,t}$, so $\eta_{s,t}^{\text{MP},k} = (1/k) \cdot (k \bar{L}^{s,t})^{-1} = 1/\sum_i L_i^{s,t}$. Recall from Section 4.3 that $\mathcal{E}(C) = \{(s, t) : s, t \notin C, s \text{ and } t \text{ connected in } G[V \setminus C]\}$ is the set of ordered pairs that the cartel C does not disconnect, and that C is called *eligible* when $\mathcal{E}(C) \neq \emptyset$. For an eligible cartel $C \subseteq V$ of size m , define the per-scheme protected-pair sets

$$\begin{aligned}\Pi_{\text{RF}}(C, \tau) &= \{(s, t) \in \mathcal{E}(C) : p_C^{s,t} \leq \tau\}, \\ \Pi_{\text{MP}}(C, k) &= \{(s, t) \in \mathcal{E}(C) : \exists i \in [k], P_i^{s,t} \cap C = \emptyset\},\end{aligned}$$

and their normalized cardinalities (protected-pair coverage fractions)

$$\begin{aligned}\pi_{\text{RF}}(C, \tau) &:= \frac{|\Pi_{\text{RF}}(C, \tau)|}{|\mathcal{E}(C)|}, & \pi_{\text{MP}}(C, k) &:= \frac{|\Pi_{\text{MP}}(C, k)|}{|\mathcal{E}(C)|}, \\ \pi_{\text{RF} \setminus \text{MP}} &:= \frac{|\Pi_{\text{RF}}(C, \tau) \setminus \Pi_{\text{MP}}(C, k)|}{|\mathcal{E}(C)|}, & \pi_{\text{MP} \setminus \text{RF}} &:= \frac{|\Pi_{\text{MP}}(C, k) \setminus \Pi_{\text{RF}}(C, \tau)|}{|\mathcal{E}(C)|}.\end{aligned}$$

Π_{RF} is the complement of the affected set $\mathcal{A}(C, \tau)$ inside $\mathcal{E}(C)$ at design threshold τ . Π_{MP} uses the path set fixed by Suurballe’s rule, so the deterministic baseline commits to its routing decision before the cartel is realized. We deliberately do not report a cartel-adaptive MP variant: identifying C at runtime falls outside both threat models considered here and outside the operational assumptions of the QKD-network layer, so the comparison of interest is between two schemes that the source can actually deploy. The two sets Π_{RF} and Π_{MP} are not nested in general. For a pair with $\kappa(s, t) = 1$ and unique cut vertex u , the deterministic baseline collapses to single-path routing ($k = 1$) through u and protects

only against cartels disjoint from that one chosen path. Random flow still distributes fragments across many sampled walks that all traverse u but otherwise visit varying intermediate nodes; for cartels C with $u \notin C$ that intersect the deterministic shortest path but miss most RF-visited nodes, (s, t) can lie in $\Pi_{\text{RF}} \setminus \Pi_{\text{MP}}$. The same effect occurs for $\kappa(s, t) \geq 2$ pairs whose two deterministic disjoint paths are both intersected by a small cartel. When two short node-disjoint paths physically avoid C , MP protects (s, t) exactly. Random flow may still fail the τ -threshold on the same (s, t) if independent walks visit C with non-negligible probability, especially on graphs where most low-degree nodes lie outside the deterministic disjoint paths. Table 10 reports, for each topology and cartel size $m \in \{1, 2, 3\}$, the topology-wide means over eligible unordered cartels of π_{RF} , π_{MP} , $\pi_{\text{RF} \setminus \text{MP}}$, and $\pi_{\text{MP} \setminus \text{RF}}$. The left $\overline{\pi_{\text{MP}}}$ column is τ -independent; the two right blocks give the RF-dependent metrics at the default $\tau = 98\%$ and at the stricter $\tau = 95\%$. Table 9 contrasts the pair-mean efficiencies $\overline{\eta}^{\lambda}$ (HS with loop-erasure, $M = 1024$) and $\overline{\eta}^{\text{MP},2}$ over biconnected ordered pairs.

Table 10. Protected-pair coverage of random flow (HS with loop-erasure, $M = 1024$) and static Suurballe $k = \kappa(s, t)$ multipath. Each cell is a mean over eligible unordered cartels C of size m . The $\overline{\pi_{\text{MP}}}$ column is independent of the RF design threshold τ ; the two right blocks report RF-dependent metrics at $\tau = 98\%$ (default) and $\tau = 95\%$.

Graph	m	$\overline{\pi_{\text{MP}}}$	$\tau = 98\%$			$\tau = 95\%$		
			$\overline{\pi_{\text{RF}}}$	$\overline{\pi_{\text{RF} \setminus \text{MP}}}$	$\overline{\pi_{\text{MP} \setminus \text{RF}}}$	$\overline{\pi_{\text{RF}}}$	$\overline{\pi_{\text{RF} \setminus \text{MP}}}$	$\overline{\pi_{\text{MP} \setminus \text{RF}}}$
NSFNET	1	100.00%	100.00%	0.00%	0.00%	100.00%	0.00%	0.00%
NSFNET	2	98.94%	100.00%	1.06%	0.00%	100.00%	1.06%	0.00%
NSFNET	3	95.55%	100.00%	4.44%	0.00%	99.62%	4.12%	0.05%
GÉANT	1	97.56%	100.00%	2.44%	0.00%	100.00%	2.44%	0.00%
GÉANT	2	94.18%	100.00%	5.82%	0.00%	99.96%	5.79%	0.01%
GÉANT	3	90.56%	99.98%	9.42%	0.00%	99.65%	9.13%	0.04%

At the default $\tau = 98\%$, random flow protects almost all eligible pairs and covers more pairs than the static Suurballe baseline in every reported row. The gap is largest on GÉANT, where $\overline{\pi_{\text{RF} \setminus \text{MP}}}$ reaches 9.42% at $m = 3$, and smaller on NSFNET, where it reaches 4.44%. Conversely, $\overline{\pi_{\text{MP} \setminus \text{RF}}}$ is 0.00% in every $\tau = 98\%$ column. The $\tau = 95\%$ block shows that the two protected-pair sets are not nested in general: $\overline{\pi_{\text{MP} \setminus \text{RF}}}$ is small but nonzero, reaching at most 0.05% on NSFNET and 0.04% on GÉANT. That asymmetry is why we treat $\tau = 98\%$ as the default operating point in the remainder of the evaluation. Figure 8 illustrates a high-impact NSFNET 3-cartel, $C = \{\text{CMI}, \text{SLC}, \text{CPK}\}$. Among all $\binom{14}{3} = 364$ unordered 3-cartels, this one has $\pi_{\text{RF} \setminus \text{MP}} = 30.91\%$ (34 of 110 eligible ordered pairs) at both $\tau = 98\%$ and $\tau = 95\%$, with $\pi_{\text{RF}} = 100.00\%$ and $\pi_{\text{MP}} = 69.09\%$. At $\tau = 95\%$, it is the maximizing 3-cartel; at $\tau = 98\%$ the maximizing 3-cartel is $\{\text{CMI}, \text{HOU}, \text{PRI}\}$, with $\pi_{\text{RF} \setminus \text{MP}} = 36.36\%$ (40 of 110 eligible ordered pairs). It is not a vertex cut, so the remaining NSFNET graph stays connected, but it hits many of the static Suurballe path sets chosen by MP. RF can still use surviving alternative walks, whereas static MP cannot change its preselected paths after the cartel is fixed.

At the default design threshold $\tau = 98\%$, the conservative design yield is $\rho^{\tau} = g_{99.99\%}^*(1024, 98\%) / 1024 = 6/1024$, supporting $\theta = 4$ after the two-fragment privacy-amplification penalty. The corresponding topology-wide means are $\overline{\eta}^{\text{RF},98} = 0.24\%$ on NSFNET and 0.11% on GÉANT. For comparison, at $\tau = 95\%$ the corresponding $\overline{\eta}^{\text{RF},95}$ values are 1.07% on NSFNET and 0.52% on GÉANT. Relative to the MP-2 column in Table 9, the $\tau = 95\%$ values are about $0.06\times$ and $0.05\times$, respectively, whereas the $\tau = 98\%$ values are only about $0.01\times$ on both topologies. The MP column in Table 9 is reported in

extracted-key bits per link-key bit ($\theta_{MP} = 1$ per execution), while the RF column uses the safe-fragment yield ρ .

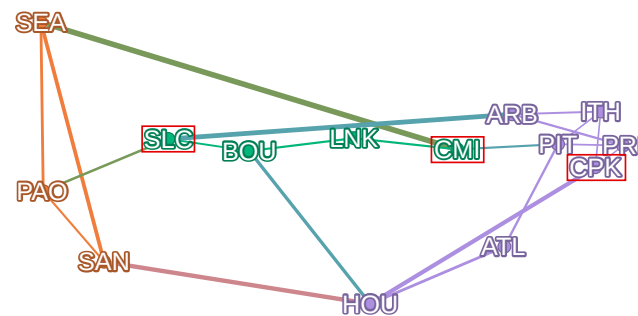


Figure 8. NSFNET with the $m = 3$ cartel $\{CMI, SLC, CPK\}$ highlighted. The highlighted node color marks the compromised cartel; other node colors are layout/community colors with no security meaning. At $\tau = 98\%$, this cartel has $\pi_{RF} = 100.00\%$, $\pi_{MP} = 69.09\%$, and $\pi_{RF \setminus MP} = 30.91\%$.

4.8. Scalability

To study scaling trends beyond the two deployed-analog topologies above, we generate a synthetic 2-vertex-connected graph with 99 nodes and 143 edges by placing clusters of 3 nodes at randomly chosen nearby coordinates and connecting each cluster to 2–3 of its 6 closest neighbors; generation is retried if the graph is not 2-vertex-connected. In Figure 9, darker nodes have higher sequence numbers (they appear later in snapshots), and edge directions indicate the insertion step when the edge was added. The graph has an average degree of ≈ 2.9 , comparable to NSFNET and GÉANT. Furthermore, we bias edge creation toward geographically closer nodes. Each node is also assigned a sequence number so we can take snapshots at $n = 3, 6, \dots, 99$ such that the snapshots maintain the same properties.

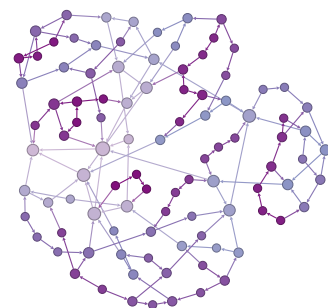


Figure 9. Synthetic graph with 99 nodes. Darker nodes have higher sequence numbers, and edge directions indicate insertion order.

Table 11 summarizes basic structural metrics of the synthetic graph.

Table 11. Topology overview of the synthetic graph. ASP is the average shortest-path length. The final columns report the percentage of ordered endpoint pairs with vertex connectivity $\kappa(s, t)$ at least 2 or 3.

Graph	Nodes	Edges	Diam.	Avg. Deg.	ASP	$\kappa \geq 2$	$\kappa \geq 3$
Generated	99	143	10	2.89	4.784	100%	21.1%

Table 12 reports single-relay exposure on the final 99-node snapshot and the peak worst-case exposure over all synthetic snapshots. To evaluate whether worst-case exposure $\max_{s,t} \chi_{s,t}$ grows with network size, we measured it on successive snapshots and report the results in Figure 10. Although the worst-case exposure remains very high throughout, we do not observe a clear monotonic increase with $|V|$. Moreover, for the non-backtracking variants, the peak maximum is not attained at the final 99-node snapshot: NB peaks at

97.3% on the 69-node snapshot, while LRV, NC, and HS peak at 97.4%, 96.5%, and 96.1%, respectively, on the 78-node snapshot. This suggests that, at least for our generated graph family, worst-case exposure is driven more by structural bottlenecks than by network size alone.

Table 12. Single-relay exposure on the synthetic graph (single-relay metric). For each variant, we report the worst-case pair exposure $\max_{s,t} \hat{\chi}$ at the final 99-node snapshot, the peak $\max_{s,t} \hat{\chi}$ over all snapshots, the snapshot size at that peak, and the median $\hat{\chi}$ over biconnected (s, t) pairs at $n = 99$.

Variant	Max $\hat{\chi}$ at $n = 99$ [%]	Peak Max $\hat{\chi}$ [%]	Peak n	Median $\hat{\chi}$ at $n = 99$ [%]
R	98.5	98.5	99	84.8
NB	97.0	97.3	69	78.5
LRV	96.4	97.4	78	77.2
NC	95.2	96.5	78	76.7
HS	93.3	96.1	78	72.8

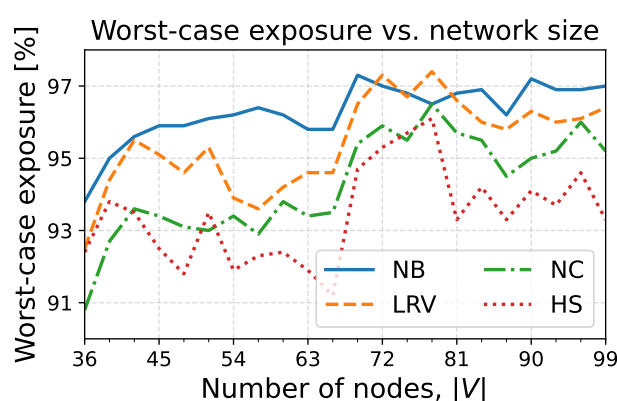


Figure 10. Worst-case exposure on synthetic graph snapshots as the network size grows.

As shown in Figure 11, the average loop-erased expected hop count over all (s, t) pairs grows approximately linearly with network size on the synthetic snapshots. Table 13 gives hop-count statistics on the final 99-node snapshot. On that graph, the average shortest-path length over 9702 ordered pairs is 4.8 (Table 11), whereas the average mean hop count is 10–18 with loop-erasure and 41–128 without it. The average length of the shortest path when original shortest-path nodes are removed is 7.5, still significantly below 10–18 post-loop-erasure. We include this metric as a proxy for the length of a second path in a node-disjoint routing scenario; it is an upper bound on the second-path length produced by Suurballe’s minimum-total-length disjoint-pair algorithm [27,28], since sequential shortest-path-then-remove need not coincide with the jointly optimal pair. Even this alternate-path baseline is shorter than the loop-erased random-flow routes on the generated graph.

Table 13. Estimated expected hop count on the final 99-node synthetic graph. Let $H_{s,t}$ be the hop count in the sampled walk from s to t . We report the average of mean, median, and maximum of $H_{s,t}$ over all biconnected (s, t) pairs, and distinguish whether loop-erasure is toggled on.

Graph	Metric	Without Loop-Erasure					With Loop-Erasure				
		R	NB	LRV	NC	HS	R	NB	LRV	NC	HS
Generated	Mean	128	68	41	42	49	10	13	16	16	18
	Median	85	46	34	35	36	9	11	14	14	16
	Max	1027	516	163	181	357	33	42	54	53	58

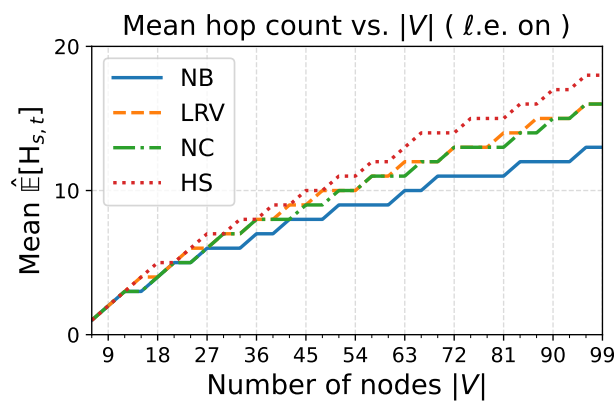


Figure 11. Average expected hop count by graph and RW variant with loop-erasure.

Tables 14 and 15 extend the efficiency and throughput analysis to the final synthetic snapshot. On the generated graph, mean raw throughput stays nearly flat across variants, at about 1.97–2.01 kbit/s, but mean model-based extracted throughput increases from about 0.41 kbit/s for NB to about 0.49 kbit/s for HS.

Table 14. Topology-wide mean RW efficiency on the final 99-node synthetic graph, in percent. Rows report $\bar{\eta}^\tau$ or $\bar{\eta}^\chi$ as in Table 6. The left block uses hop counts from the original walk, while the right block uses hop counts after loop-erasure. In both blocks, exposure is measured on the original walk and $M = 1024$.

Graph	Mean	Loop-Erasure off [%]				Loop-Erasure on [%]			
		NB	LRV	NC	HS	NB	LRV	NC	HS
Generated	$\bar{\eta}^\tau$	0.05	0.06	0.09	0.10	0.12	0.09	0.15	0.16
	$\bar{\eta}^\chi$	3.5	3.7	3.7	3.7	4.6	4.4	4.4	4.4

Table 15. Mean throughput on the final 99-node synthetic graph. The top block reports raw fragment throughput $T_{s,t}$ in kbit/s. The bottom block reports mean model-based extracted throughput $R_{s,t} = T_{s,t}\rho_{s,t}$ in kbit/s under the simplified entropy model.

Metric	Graph	NB	LRV	NC	HS
$T_{s,t}$	Generated	1.97	1.97	1.97	2.01
$R_{s,t}$	Generated	0.41	0.43	0.44	0.49

4.9. Evaluation Summary

The numerical results above support five main conclusions. First, HS with loop-erasure is the only evaluated random-flow variant that we would carry forward as a protocol candidate. It gives the lowest worst-case exposure among the evaluated variants, while NC is only a diagnostic comparison because it assumes global node-set knowledge. R, NB, and LRV are therefore best read as baselines rather than implementer-facing recommendations. Second, loop-erasure is what makes the random-flow route lengths plausible. On GÉANT, the mean HS route length falls from 19.8 hops to 8.7 hops after loop-erasure; on the synthetic graph, it falls from 49 hops to 18 hops. The same pattern holds for the other nontrivial random-walk variants. Nevertheless, these loop-erased routes remain longer than shortest-path and alternate-path proxies for topology-aware node-disjoint routing [27,28], so random flow should not be interpreted as a route-efficiency improvement over a controller that already knows suitable disjoint paths. Third, the efficiency and throughput results separate pair-specific behavior from conservative worst-case parameterization. With pair-specific exposure estimates, the model-based extracted-throughput proxy is tightly clustered on NSFNET, at 0.62–0.64 kbit/s, while on GÉANT it increases

from 0.50 kbit/s for NB to 0.56 kbit/s for HS. With the default $\tau = 98\%$ design threshold, however, the conservative RF efficiency falls to 0.24% on NSFNET and 0.11% on GÉANT, about $0.01 \times$ the MP-2 baseline in both cases. The method is therefore attractive only when the deployment accepts this worst-case cost, uses pair-specific parameterization, or treats random flow as a topology-oblivious fallback rather than as the efficiency optimum. Fourth, the comparison with static multipath is mixed rather than one-sided. Static Suurballe-style multipath remains the natural benchmark when route computation and topology knowledge are acceptable. At the same time, at $\tau = 98\%$, random flow protects almost all eligible pairs and covers more pairs than the static multipath baseline in every reported protected-pair row; on GÉANT with $m = 3$, the mean RF-only coverage reaches 9.42% and the mean MP-only coverage is 0.00%. At $\tau = 95\%$, the protected-pair sets are not nested, so neither method dominates the other for every cartel placement. Fifth, the two- and three-relay GÉANT cartel analysis is a sensitivity study, not a full multi-relay ITS guarantee. For arbitrary eligible cartels, the median affected share is often zero at the higher design thresholds, but the maximum column shows that favorable cartel placements can still make the single-relay design exposure too optimistic: at $\tau = 98\%$, the worst three-node cartel affects 19.49% of eligible ordered pairs, whereas arbitrary three-node placements affect only 0.02% on average. When this happens, the remaining claim is the computational fallback of Section 3.5, not the information-theoretic extraction claim.

5. Conclusions

This paper presents random flow as a topology-oblivious QKD key-relaying design that trades extracted-key efficiency for local forwarding, simpler operation, and reduced dependence on precomputed disjoint paths. Random flow fragments key material, sends fragments along independently sampled random walks, and applies seeded privacy amplification at the destination. The evaluation identifies HS forwarding with scouting-based loop erasure as the useful protocol candidate. On GÉANT, HS reduces worst-case single-relay exposure to 91.3%, compared with 98.3% for simple random walk, and loop erasure reduces mean HS route length from 19.8 to 8.7 hops. The main observation is that a simple local construction can exploit partial disjointness: although static node-disjoint multipath remains the natural benchmark when topology knowledge and route computation are acceptable, random flow achieves competitive protected-pair coverage under the evaluated cartel metric. The price is efficiency under conservative worst-case parameterization. At $\tau = 98\%$, RF efficiency is low compared with MP-2, but the throughput simulations suggest that this overhead is partly mitigated in the two-KME communication scenario because random walks spread fragment load over the network rather than concentrating it on a fixed route set. The computational fallback in Section 3.5 further explains why a large M remains useful even when the ITS extracted block is small. Random flow therefore points to a different design target for Q-KMS systems: not maximum route efficiency, but simple, decentralized, auditable key relaying with minimal routing state. This is particularly relevant in deployments where global topology knowledge, centralized orchestration, or per-pair path computation are operational liabilities. Several extensions could improve scalability without abandoning this design direction. Signed forwarding histories would make realized-path accounting auditable. Scouting can be extended so that the source samples several candidate walks and relays payload only over a selected subset, for example the most mutually disjoint scouts. Larger networks could also be handled hierarchically by dividing the topology into modularity communities and applying random flow within or between those regions.

Author Contributions: K.P. and S.K. developed the research concept and methodology. K.P. implemented the software, conducted the main evaluation, performed the analysis, and prepared the original

draft. R.I. contributed to numerical evaluation and validation. J.V. provided supervision and project leadership. E.K., E.C., L.L., and E.R. contributed through discussion, review, and support within the broader project. All authors have read and agreed to the published version of the manuscript.

Funding: This work has been supported by the Latvian Quantum Initiative under European Union Recovery and Resilience Facility project no. 2.3.1.1.i.0/1/22/I/CFLA/001.

Data Availability Statement: The graph edge lists and simulation source code supporting the findings of this study are publicly available at <https://github.com/LUMII-Syslab/random-walk-key-relaying> (accessed on 4 February 2026).

Conflicts of Interest: The authors declare no conflicts of interest.

References

- Shannon, C.E. Communication Theory of Secrecy Systems. *Bell Syst. Tech. J.* **1949**, *28*, 656–715. [\[CrossRef\]](#)
- Bennett, C.H.; Brassard, G. Quantum Cryptography: Public Key Distribution and Coin Tossing. *Theor. Comput. Sci.* **2014**, *560*, 7–11. [\[CrossRef\]](#)
- Shor, P.W.; Preskill, J. Simple Proof of Security of the BB84 Quantum Key Distribution Protocol. *Phys. Rev. Lett.* **2000**, *85*, 441–444. [\[CrossRef\]](#) [\[PubMed\]](#)
- Peev, M.; Pacher, C.; Alléaume, R.; Barreiro, C.; Bouda, J.; Boxleitner, W.; Debuisschert, T.; Diamanti, E.; Dianati, M.; Dynes, J.F.; et al. The SECOQC Quantum Key Distribution Network in Vienna. *New J. Phys.* **2009**, *11*, 075001. [\[CrossRef\]](#)
- Stucki, D.; Legré, M.; Buntschu, F.; Clausen, B.; Felber, N.; Gisin, N.; Henzen, L.; Junod, P.; Litzistorf, G.; Monbaron, P.; et al. Long-Term Performance of the SwissQuantum Quantum Key Distribution Network in a Field Environment. *New J. Phys.* **2011**, *13*, 123001. [\[CrossRef\]](#)
- Chen, T.Y.; Jiang, X.; Tang, S.B.; Zhou, L.; Yuan, X.; Zhou, H.; Wang, J.; Liu, Y.; Chen, L.K.; Liu, W.Y.; et al. Implementation of a 46-Node Quantum Metropolitan Area Network. *npj Quantum Inf.* **2021**, *7*, 134. [\[CrossRef\]](#)
- Gisin, N.; Ribordy, G.; Tittel, W.; Zbinden, H. Quantum Cryptography. *Rev. Mod. Phys.* **2002**, *74*, 145–195. [\[CrossRef\]](#)
- Huang, D.; Huang, P.; Lin, D.; Zeng, G. Long-Distance Continuous-Variable Quantum Key Distribution by Controlling Excess Noise. *Sci. Rep.* **2016**, *6*, 19201. [\[CrossRef\]](#) [\[PubMed\]](#)
- Pittaluga, M.; Lo, Y.S.; Brzosko, A.; Woodward, R.I.; Scalcon, D.; Winnel, M.S.; Roger, T.; Dynes, J.F.; Owen, K.A.; Juárez, S.; et al. Long-Distance Coherent Quantum Communications in Deployed Telecom Networks. *Nature* **2025**, *640*, 911–917. [\[CrossRef\]](#) [\[PubMed\]](#)
- Salvail, L.; Peev, M.; Diamanti, E.; Alléaume, R.; Lütkenhaus, N.; Länger, T. Security of Trusted Repeater Quantum Key Distribution Networks. *J. Comput. Secur.* **2010**, *18*, 61–87. [\[CrossRef\]](#)
- Dervisevic, E.; Tankovic, A.; Fazel, E.; Kompella, R.; Fazio, P.; Voznak, M.; Mehic, M. Quantum Key Distribution Networks-Key Management: A Survey. *ACM Comput. Surv.* **2025**, *57*, 257. [\[CrossRef\]](#)
- Cvitić, I.; Peraković, D.; Nolasco, A. Overview of Routing Approaches in Quantum Key Distribution Networks. *arXiv* **2025**, arXiv:2511.15465.
- Moy, J. RFC 2328: OSPF Version 2. 1998. Available online: <https://www.rfc-editor.org/info/rfc2328/> (accessed on 4 February 2026).
- Drif, Y.; Bedhief, I.; Chatzinotas, S. Distributed Key Relay: OSPF for Effective QKD. *IEEE Commun. Stand. Mag.* **2025**, *10*, 154–161. [\[CrossRef\]](#)
- Liao, S.K.; Cai, W.Q.; Liu, W.Y.; Zhang, L.; Li, Y.; Ren, J.G.; Yin, J.; Shen, Q.; Cao, Y.; Li, Z.P.; et al. Satellite-to-Ground Quantum Key Distribution. *Nature* **2017**, *549*, 43–47. [\[CrossRef\]](#)
- Pirandola, S.; Laurenza, R.; Ottaviani, C.; Banchi, L. Fundamental Limits of Repeaterless Quantum Communications. *Nat. Commun.* **2017**, *8*, 15043. [\[CrossRef\]](#) [\[PubMed\]](#)
- Lucamarini, M.; Yuan, Z.L.; Dynes, J.F.; Shields, A.J. Overcoming the Rate–Distance Limit of Quantum Key Distribution without Quantum Repeaters. *Nature* **2018**, *557*, 400–403. [\[CrossRef\]](#) [\[PubMed\]](#)
- Minder, M.; Pittaluga, M.; Roberts, G.L.; Lucamarini, M.; Dynes, J.F.; Yuan, Z.L.; Shields, A.J. Experimental Quantum Key Distribution beyond the Repeaterless Secret Key Capacity. *Nat. Photonics* **2019**, *13*, 334–338. [\[CrossRef\]](#)
- Elliott, C. Building the Quantum Network. *New J. Phys.* **2002**, *4*, 46–46. [\[CrossRef\]](#)
- ETSI GS QKD 014 V1.1.1; Quantum Key Distribution (QKD); Protocol and Data Format of REST-Based Key Delivery API. European Telecommunications Standards Institute (ETSI): Sophia Antipolis Cedex, France, 2019.
- Rescorla, E. RFC 8446: The Transport Layer Security (TLS) Protocol Version 1.3. 2018. Available online: <https://www.rfc-editor.org/info/rfc8446/> (accessed on 4 February 2026).
- Kaufman, C.; Hoffman, P.; Nir, Y.; Eronen, P.; Kivinen, T. RFC 7296: Internet Key Exchange Protocol Version 2 (IKEv2). 2014. Available online: <https://www.rfc-editor.org/info/rfc7296/> (accessed on 4 February 2026).

23. Dowling, B.; Hansen, T.B.; Paterson, K.G. Many a Mickle Makes a Muckle: A Framework for Provably Quantum-Secure Hybrid Key Exchange. In *Post-Quantum Cryptography*; Ding, J., Tillich, J.P., Eds.; Springer International Publishing: Berlin/Heidelberg, Germany, 2020; Volume 12100, pp. 483–502. [\[CrossRef\]](#)
24. Beals, T.R.; Sanders, B.C. Distributed Relay Protocol for Probabilistic Information-Theoretic Security in a Randomly-Compromised Network. In *Proceedings of the International Conference on Information Theoretic Security*; Springer: Berlin/Heidelberg, Germany, 2008. [\[CrossRef\]](#)
25. Wen, H.; Han, Z.; Zhao, Y.; Guo, G.; Hong, P. Multiple Stochastic Paths Scheme on Partially-Trusted Relay Quantum Key Distribution Network. *Sci. China Ser. Inf. Sci.* **2009**, *52*, 18–22. [\[CrossRef\]](#)
26. Kiktenko, E.O.; Tayduganov, A.; Fedorov, A.K. Routing Algorithm Within the Multiple Non-Overlapping Paths' Approach for Quantum Key Distribution Networks. *Entropy* **2024**, *26*, 1102. [\[CrossRef\]](#) [\[PubMed\]](#)
27. Suurballe, J.W. Disjoint Paths in a Network. *Networks* **1974**, *4*, 125–145. [\[CrossRef\]](#)
28. Bhandari, R. *Survivable Networks: Algorithms for Diverse Routing*; Kluwer Academic Publishers: Boston, MA, USA, 1998.
29. Kumar, P.; Kundu, N.K.; Kar, B. Quantum Key Distribution Routing Protocol in Quantum Networks: Overview and Challenges. *arXiv* **2024**, arXiv:2407.13156. [\[CrossRef\]](#)
30. Yao, J.; Wang, Y.; Li, Q.; Mao, H.; El-Latif, A.A.A.; Chen, N. An Efficient Routing Protocol for Quantum Key Distribution Networks. *Entropy* **2022**, *24*, 911. [\[CrossRef\]](#) [\[PubMed\]](#)
31. Álvarez Roa, M.; Stan, C.; Verschoor, S.; Tafur Monroy, I.; Rommel, S. Decentralized Key Distribution versus On-Demand Relaying for QKD Networks. *J. Opt. Commun. Netw.* **2025**, *17*, 732. [\[CrossRef\]](#)
32. Stan, C.; Verchere, D.; Olmos, J.J.V.; Tafur Monroy, I.; Rommel, S. Dynamic-Threshold-Based Pre-Relaying for Enhanced Key Allocation in Quantum-Secured Networks. *J. Opt. Commun. Netw.* **2025**, *17*, 233. [\[CrossRef\]](#)
33. Wang, M.; Li, J.; Xue, K.; Li, R.; Yu, N.; Li, Y.; Liu, Y.; Sun, Q.; Lu, J. A Segment-Based Multipath Distribution Method in Partially-Trusted Relay Quantum Networks. *IEEE Commun. Mag.* **2023**, *61*, 184–190. [\[CrossRef\]](#)
34. Le, Q.C.; Bellot, P.; Demaille, A. Towards the World-Wide Quantum Network. In *Information Security Practice and Experience*; Chen, L., Mu, Y., Susilo, W., Eds.; Springer: Berlin/Heidelberg, Germany, 2008; Volume 4991, pp. 218–232. [\[CrossRef\]](#)
35. Le Quoc, C.; Bellot, P.; Demaille, A. Stochastic Routing in Large Grid-Shaped Quantum Networks. In *Proceedings of the 2007 IEEE International Conference on Research, Innovation and Vision for the Future*; IEEE: Piscataway, NJ, USA, 2007; pp. 166–174. [\[CrossRef\]](#)
36. Ghourab, E.M.; Azab, M.; Gračanin, D. QKD Routing Scheme for a Zero-Trust Network System: A Moving Target Defense Approach. *Big Data Cogn. Comput.* **2025**, *9*, 76. [\[CrossRef\]](#)
37. Braunstein, S.L.; Pirandola, S. Side-Channel-Free Quantum Key Distribution. *Phys. Rev. Lett.* **2012**, *108*, 130502. [\[CrossRef\]](#)
38. Lo, H.K.; Curty, M.; Qi, B. Measurement-Device-Independent Quantum Key Distribution. *Phys. Rev. Lett.* **2012**, *108*, 130503. [\[CrossRef\]](#) [\[PubMed\]](#)
39. Liu, Y.; Chen, T.Y.; Wang, L.J.; Liang, H.; Shentu, G.L.; Wang, J.; Cui, K.; Yin, H.L.; Liu, N.L.; Li, L.; et al. Experimental Measurement-Device-Independent Quantum Key Distribution. *Phys. Rev. Lett.* **2013**, *111*, 130502. [\[CrossRef\]](#) [\[PubMed\]](#)
40. Fan-Yuan, G.J.; Lu, F.Y.; Wang, S.; Yin, Z.Q.; He, D.Y.; Chen, W.; Zhou, Z.; Wang, Z.H.; Teng, J.; Guo, G.C.; et al. Robust and Adaptable Quantum Key Distribution Network without Trusted Nodes. *Optica* **2022**, *9*, 812–823. [\[CrossRef\]](#)
41. Yan, W.; Zheng, X.; Wen, W.; Lu, L.; Du, Y.; Lu, Y.Q.; Zhu, S.; Ma, X.S. A Measurement-Device-Independent Quantum Key Distribution Network Using Optical Frequency Comb. *npj Quantum Inf.* **2025**, *11*, 97. [\[CrossRef\]](#)
42. Cao, Y.; Zhao, Y.; Li, J.; Lin, R.; Zhang, J.; Chen, J. Hybrid Trusted/Untrusted Relay-Based Quantum Key Distribution Over Optical Backbone Networks. *IEEE J. Sel. Areas Commun.* **2021**, *39*, 2701–2718. [\[CrossRef\]](#)
43. Yu, X.; Liu, Y.; Zou, X.; Cao, Y.; Zhao, Y.; Nag, A.; Zhang, J. Secret-Key Provisioning With Collaborative Routing in Partially-Trusted-Relay-based Quantum-Key-Distribution-Secured Optical Networks. *J. Light. Technol.* **2022**, *40*, 3530–3545. [\[CrossRef\]](#)
44. Rass, S.; Mehic, M.; Voznak, M.; König, S. Hacking the Least Trusted Node: Indirect Eavesdropping in Quantum Networks. *IEEE Access* **2024**, *12*, 160973–160981. [\[CrossRef\]](#)
45. Lydersen, L.; Wiechers, C.; Wittmann, C.; Elser, D.; Skaar, J.; Makarov, V. Hacking Commercial Quantum Cryptography Systems by Tailored Bright Illumination. *Nat. Photonics* **2010**, *4*, 686–689. [\[CrossRef\]](#)
46. Ford, L.R.; Fulkerson, D.R. *Flows in Networks*; Princeton University Press: Princeton, NJ, USA, 1962.
47. Alon, N.; Benjamini, I.; Lubetzky, E.; Sodin, S. Non-Backtracking Random Walks Mix Faster. *Commun. Contemp. Math.* **2007**, *9*, 585–603. [\[CrossRef\]](#)
48. Akash, A.K.; Fekete, S.; Lee, S.K.; López-Ortiz, A.; Maftuleac, D.; McLurkin, J. Lower Bounds for Graph Exploration Using Local Policies. *J. Graph Algorithms Appl.* **2017**, *21*, 371–387. [\[CrossRef\]](#)
49. Grover, L.K. A Fast Quantum Mechanical Algorithm for Database Search. In *Proceedings of the Twenty-Eighth Annual ACM Symposium on Theory of Computing-STOC '96*; ACM Press: New York, NY, USA, 1996; pp. 212–219. [\[CrossRef\]](#)
50. Fung, C.H.F.; Ma, X.; Chau, H.F. Practical Issues in Quantum-Key-Distribution Postprocessing. *Phys. Rev. A—At. Mol. Opt. Phys.* **2010**, *81*, 012318. [\[CrossRef\]](#)

51. Bennett, C.H.; Brassard, G.; Robert, J.M. Privacy Amplification by Public Discussion. *SIAM J. Comput.* **1988**, *17*, 210–229. [\[CrossRef\]](#)
52. Renner, R.; König, R. Universally Composable Privacy Amplification Against Quantum Adversaries. In *Theory of Cryptography*; Kilian, J., Ed.; Springer: Berlin/Heidelberg, Germany, 2005; Volume 3378, pp. 407–425. [\[CrossRef\]](#)
53. Tomamichel, M.; Schaffner, C.; Smith, A.; Renner, R. Leftover Hashing Against Quantum Side Information. *IEEE Trans. Inf. Theory* **2011**, *57*, 5524–5535. [\[CrossRef\]](#)
54. Stinson, D.R.; Paterson, M.B. *Cryptography: Theory and Practice*, 4th ed.; Chapman and Hall/CRC: Boca Raton, FL, USA, 2005. [\[CrossRef\]](#)
55. Lawler, G.F.; Limic, V. *Random Walk: A Modern Introduction*, 1st ed.; Cambridge University Press: New York, NY, USA, 2010. [\[CrossRef\]](#) [\[PubMed\]](#)
56. Wilson, D.B. Generating Random Spanning Trees More Quickly than the Cover Time. In *Proceedings of the Twenty-Eighth Annual ACM Symposium on Theory of Computing-STOC '96*; ACM Press: New York, NY, USA, 1996; pp. 296–303. [\[CrossRef\]](#)
57. Mehic, M.; Niemiec, M.; Rass, S.; Ma, J.; Peev, M.; Aguado, A.; Martin, V.; Schauer, S.; Poppe, A.; Pacher, C.; et al. Quantum Key Distribution: A Networking Perspective. *ACM Comput. Surv. (CSUR)* **2020**, *53*, 1–41. [\[CrossRef\]](#)
58. Mukherjee, B.; Ramamurthy, S.; Banerjee, D.; Mukherjee, A. Some Principles for Designing a Wide-Area Optical Network. In *Proceedings of the INFOCOM '94 Conference on Computer Communications*; IEEE: Piscataway, NJ, USA, 2002; pp. 110–119. [\[CrossRef\]](#)
59. GÉANT. GN4-3N. Web Page Describing the GN4-3N Project (2019–2023); Topology Diagram Consulted to Reconstruct Node/Edge Lists (Links > 1000 km Removed in Our Processed Graph). 2023. Available online: <https://network.geant.org/gn4-3n/> (accessed on 4 February 2026).
60. Xuereb, A. Towards an Ultra-Secure Communication Network for the EU. GÉANT CONNECT. Online Article Discussing EuroQCI-Oriented Ultra-Secure Networking Efforts Including QKD. 2023. Available online: <https://connect.geant.org/2023/12/13/towards-an-ultra-secure-communication-network-for-the-eu> (accessed on 4 February 2026).

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.