

Experimental demonstration of entanglement delivery using a quantum network stack

Pompili, M.; Delle Donne, C.; te Raa, I.; van der Vecht, B.; Skrzypczyk, M.; Ferreira, G.; de Kluijver, L.; Stolk, A. J.; Hermans, S. L.N.; Pawełczak, P.

DOI

[10.1038/s41534-022-00631-2](https://doi.org/10.1038/s41534-022-00631-2)

Publication date

2022

Document Version

Final published version

Published in

NPJ Quantum Information

Citation (APA)

Pompili, M., Delle Donne, C., te Raa, I., van der Vecht, B., Skrzypczyk, M., Ferreira, G., de Kluijver, L., Stolk, A. J., Hermans, S. L. N., Pawełczak, P., Kozłowski, W., Hanson, R., & Wehner, S. (2022). Experimental demonstration of entanglement delivery using a quantum network stack. *NPJ Quantum Information*, 8(1), 10. Article 121. <https://doi.org/10.1038/s41534-022-00631-2>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

ARTICLE OPEN



Experimental demonstration of entanglement delivery using a quantum network stack

M. Pompili^{1,2}, C. Delle Donne^{1,2}, I. te Raa¹, B. van der Vecht¹, M. Skrzypczyk¹, G. Ferreira¹, L. de Kluijver¹, A. J. Stolk¹, S. L. N. Hermans¹, P. Pawełczak¹, W. Kozłowski¹, R. Hanson¹✉ and S. Wehner¹✉

Scaling current quantum communication demonstrations to a large-scale quantum network will require not only advancements in quantum hardware capabilities, but also robust control of such devices to bridge the gap in user demand. Moreover, the abstraction of tasks and services offered by the quantum network should enable platform-independent applications to be executed without the knowledge of the underlying physical implementation. Here we experimentally demonstrate, using remote solid-state quantum network nodes, a link layer, and a physical layer protocol for entanglement-based quantum networks. The link layer abstracts the physical-layer entanglement attempts into a robust, platform-independent entanglement delivery service. The system is used to run full state tomography of the delivered entangled states, as well as preparation of a remote qubit state on a server by its client. Our results mark a clear transition from physics experiments to quantum communication systems, which will enable the development and testing of components of future quantum networks.

npj Quantum Information (2022)8:121; <https://doi.org/10.1038/s41534-022-00631-2>

INTRODUCTION

By sharing entangled states over large distances, the future Quantum Internet^{1,2} can unlock new possibilities in secure communication³, distributed and blind quantum computation^{4,5}, and metrology^{6,7}. Fundamental primitives for entanglement-based quantum networks have been demonstrated across several physical platforms, including trapped ions^{8,9}, neutral atoms^{10,11}, diamond color centers^{12–15}, and quantum dots^{16,17}. To scale up such physics experiments to intermediate-scale quantum networks, researchers have been investigating how to enclose the complex nature of quantum entanglement generation into more robust abstractions^{18–24}.

A common way to facilitate the scalability of complex systems is to break down their architecture into a stack of layers. Each layer in such a stack is characterized by a specific service that it provides to the layers above, reducing complexity for the higher layers, which can subsequently rely on this service. Moreover, the higher layers need no knowledge of the specific protocol and physical realization that a lower layer uses to realize the specified service. An example from classical networking is the TCP/IP stack used on the present-day Internet. In this stack, the link layer enables reliable transmission of data between two network nodes that are directly connected by an unreliable physical medium such as fiber or radio. Higher layers can rely on errors being detected by the link layer and are agnostic about whether the underlying link layer protocol is Ethernet or Wi-Fi.

Several network stacks have been proposed for quantum network nodes^{19–21}, like the one depicted in Fig. 1. These draw inspiration from classical architectures like the TCP/IP stack or the more generic Open Systems Interconnection (OSI) model. Specifically, the functional allocation of the stack proposed in ref. 19 conceptually mirrors the TCP/IP stack in that the link layer ensures reliable (quantum) communication between adjacent nodes, and the network layer extends this service to nodes not directly connected by a physical medium themselves. We

emphasize that, of course, no quantum data is passed up and down the layers of the stack, but only qubit metadata. Very intuitively, such metadata is similar to passing only references to an address in a physical memory up and down the stack (similar to what happens in many implementations of the TCP/IP stack in practice), while in the classical case, data may of course also be copied up and down layers.

We also note that the Quantum Internet and the associated quantum network stack, do not aim to replace the classical Internet—they will likely coexist, as the Quantum Internet cannot operate without classical communication in practice. In addition to classical information used to facilitate entanglement generation, we also expect classical communication at the level of the quantum application itself (e.g., quantum key distribution), which would, for practical reasons, be performed using the classical Internet. Finally, in a quantum network, classical communication could also be used to realize controllers like those at the core of software-defined networks (SDN)²⁵ to distribute information for resource scheduling and quality of service²⁶. The proposed quantum network stack architecture, along with proposals for resource scheduling and routing techniques (e.g., refs. 26–33), pave the way for larger-scale quantum networks.

In this work, we experimentally demonstrate a link layer protocol for entanglement-based quantum networks. The link layer abstracts the generation of entangled states between two physically separated solid-state qubits into a robust and platform-independent service. An application can request entangled states from the link layer and then, in addition, apply local quantum operations on the entangled qubits in real-time. Using the link layer, we perform full state tomography of the generated states and achieve remote state preparation—a building block for blind quantum computation—as well as measure the latency of the entanglement generation service.

To evaluate the correct operation and performance of our system, we measure (a) the fidelity of the generated states and (b)

¹QuTech & Kavli Institute of Nanoscience, Delft University of Technology, 2628 CJ Delft, The Netherlands. ²These authors contributed equally: M. Pompili and C. Delle Donne. ✉email: R.Hanson@tudelft.nl; S.D.C.Wehner@tudelft.nl

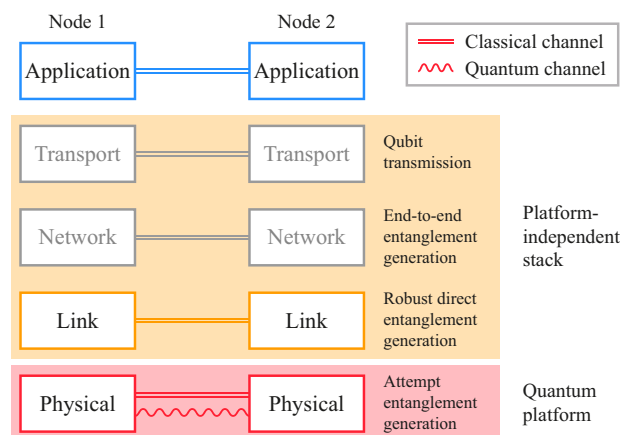


Fig. 1 Quantum network stack architecture. At the bottom of the stack, the physical layer (red), which is highly quantum-platform-dependent, is tasked with attempting entanglement generation. The link layer (yellow) uses the functionality provided by the physical layer to provide a platform-independent and robust entanglement generation service between neighboring nodes to the higher layers. The network and transport layer (not implemented in this work, grayed out) will support end-to-end connectivity and qubit transmission. Applications (blue) use the services offered by the stack to perform quantum networking tasks. Based on Dahlberg et al.¹⁹.

the latency incurred by the link layer and physical layer when generating entangled pairs. For both fidelity and latency, we find that our system performs with a marginal overhead with respect to previous non-platform-independent experiments. We also identify the sources of the additional overhead incurred and propose improvements for future realizations.

RESULTS

Quantum link layer protocol

Remote entanglement generation constitutes a fundamental building block of quantum networking. However, for a user to be able to integrate it into more complex quantum networking applications and protocols, the entanglement generation service must also be: (a) robust, meaning that the user should not have to deal with entanglement failures and retries, and that an entanglement request should result in the delivery of an entangled pair; (b) quantum-platform-independent, in order for the user to be able to request entanglement without having to understand the inner workings of the underlying physical implementation; (c) on-demand, such that the user can request and consume entanglement as part of a larger quantum communication application. Robust, platform-independent, on-demand entanglement generation must figure as one of the basic services offered by a system running on a quantum network node. In other words, establishing a reliable quantum link between two directly connected nodes is the task of the first layer above the physical layer in a quantum networking protocol stack, as portrayed in Fig. 1. Following the TCP/IP stack nomenclature, we refer to this layer as the link layer. We remark that, in the framework of a multi-node network, a quantum network stack should also feature a network layer (called internet layer in the TCP/IP model) to establish links between non-adjacent nodes and, optionally, a transport layer to encapsulate qubit transmission into a service^{19–21} (as shown in Fig. 1).

Link layer service. The service provided by a link layer protocol for quantum networks should expose a few configuration parameters to its user. To ensure a platform-independent

interaction with the link layer, such parameters should be common to all possible implementations of the quantum physical device. In this work, we implement a revised version of the link layer protocol proposed—but not implemented—in ref. ¹⁹, with the following service description. The interface exposed by the link layer should allow the higher layer to specify: (a) Remote node ID, an identifier of the remote node to produce entanglement with (in case the requesting node has multiple neighbors); (b) Number of entangled pairs, to allow for the creation of several pairs with one request; (c) Minimum fidelity, an indication of the desired minimum fidelity for the produced pairs; (d) Delivery type, whether to keep the produced pair for future use (type K), measure it directly after creation (type M), or measure the local qubit immediately and instruct the remote node to keep its own for future use (type R, used for remote state preparation); (e) Measurement basis, the basis to use when measuring M- or R-type entangled pairs; (f) Request timeout, to indicate a time limit for the processing of the request. After submitting an entanglement generation request, the user should expect the link layer to coordinate with the remote node and to handle entanglement generation attempts and retries until all the desired pairs are produced (or until the timeout has expired). When completing an entanglement generation request, the link layer should then report to the above layer the following: (a) Produced Bell state, the result of entanglement generation; (b) Measurement outcome, in case of M- or R-type entanglement requests; (c) Entanglement ID, to uniquely identify an entangled pair consistently across source and destination of the request.

Quantum link layer protocol. A design of a quantum link layer protocol that offers the above service is the quantum entanglement generation protocol (QEGP) proposed by Dahlberg et al.¹⁹. As originally designed, this protocol relies on the underlying quantum physical layer protocol to achieve accurate timing synchronization with its remote peer and to detect inconsistencies between the local state and the state of the remote counterpart. To satisfy such requirements, QEGP is accompanied by a quantum physical layer protocol, called midpoint heralding protocol (MHP), designed to support QEGP on heralded entanglement-based quantum links.

Entanglement requests and agreement. QEGP exposes an interface for its user to submit entanglement requests. An entanglement request can specify all the aforementioned configuration parameters (remote node ID, number of entangled pairs, minimum fidelity, request type, measurement basis), and an additional set of parameters which can be used to determine the priority of the request. In the theoretical protocol proposed in ref. ¹⁹, agreement on the requests between the nodes is achieved using a distributed queue protocol (DQP) which adds the incoming requests to a joint queue. The distributed queue, managed by the node designated as primary, ensures that both nodes schedule pending entanglement requests in the same order. Moreover, QEGP attaches a timestamp to each request in the distributed queue, so that both nodes can process the same entanglement request simultaneously.

Time synchronization. Time-scheduling entanglement generation requests is necessary for the two neighboring nodes to trigger entanglement generation at the same time, and avoid wasting entanglement attempts. QEGP relies on MHP to maintain and distribute a synchronized clock, which QEGP itself uses to schedule entanglement requests. The granularity of such a clock is only marginally important, but its consistency across the two neighboring nodes is paramount to make sure that entanglement attempts are triggered simultaneously on the two ends.

Mismatch verification. One of the main responsibilities of MHP is to verify that both nodes involved in entanglement generation are servicing the same QEGP request at the same time, which the protocol achieves by sending an auxiliary classical message to the heralding station when the physical device sends the flying-qubit. The heralding station can thus verify that the messages fetched by the two MHP peers are consistent and correspond to the same QEGP request.

QEGP challenges. We identify three main challenges that would be faced when deploying QEGP on a large-scale quantum network, while suggesting an alternative solution for each of these. (C1) Using a link-local protocol (DQP) to schedule entanglement requests, albeit sufficient for a single-link network, becomes challenging in larger networks, given that a node might be connected to more than just one peer. In such scenarios, the scheduling of entanglement requests can instead be deferred to a centralized scheduling entity, one which has more comprehensive knowledge of the entire (sub)network²⁶. (C2) Entrusting the triggering of entanglement attempts to QEGP would impose very stringent real-time constraints on the system where QEGP itself is deployed—even microsecond-level latencies on either side of the link can result in out-of-sync (thus wasteful) entanglement attempts. While ref. ¹⁹ identifies this problem as well, the original MHP protocol assumes that both QEGP peers issue an entanglement command to the physical layer at the same clock cycle. In this scheme, MHP initiates an entanglement attempt regardless of the state of the remote counterpart. We believe that fine-grained entanglement attempt synchronization should pertain to the physical layer only, building on the assumption that the real-time controllers deployed at the physical layer of each node are anyway highly synchronized¹⁵. (C3) Checking for request mismatches at the heralding station requires the latter to be capable of performing such checks in real-time. Given that the two neighboring MHP protocols have to anyway synchronize before attempting entanglement, we suggest that, as an alternative approach, consistency checks be performed at the nodes themselves, rather than at the heralding station, just before entering the entanglement attempt routine.

Revised protocol

To address the present QEGP and MHP challenges with the proposed solutions, we have made some modifications to the original design of the two protocols. In particular, we adopted a centralized request scheduling mechanism²⁶ to tackle challenge (C1), we delegated the ultimate triggering of entanglement attempts to MHP as a solution to challenge (C2), and we assigned request mismatch verification to the MHP protocol running on each node, rather than to the heralding station, to address the challenge (C3).

Centralized request scheduling. To avoid using a link-local protocol (DQP) to schedule entanglement requests, our version of QEGP defers request scheduling to a centralized request scheduler, whereby a node's entanglement generation schedule is computed on the basis of the whole network's needs. Delegating network scheduling jobs to centralized entities is, albeit not the only alternative, a common paradigm of classical networks, and especially of software-defined networking (SDN)—a concept that has been recently investigated in the context of quantum networking^{22,23}. In large networks, such controllers are logically centralized, but physically distributed to ensure their reliability and availability in spite of possible failures. In our system, the centralized scheduler produces a time-division multiple access (TDMA) network schedule—one for each node in the network—where each time bin is reserved for a certain class of entanglement generation requests²⁶. A class of requests may comprise, for

instance, all requests coming from the same application and asking for the same fidelity of the entangled states. While reserving time bins may be redundant in a single-link network, integrating a centralized scheduling mechanism early on into the link layer protocol will facilitate future developments.

MHP synchronization and timeout. Although centralized request scheduling makes the synchronization of QEGP peers easier, precise triggering of entanglement attempts should still be entrusted to the component of the system where time is the most deterministic—in our case, the physical layer protocol MHP. In contrast to ref. ¹⁹, once MHP fetches an entanglement instruction from QEGP, the protocol announces itself as ready to its remote peer, and waits for the latter to do so as well. After this synchronization step succeeds, the two MHP peers can instruct the underlying hardware to trigger an entanglement attempt at a precise point in time. If instead, one of the two MHP peers does not receive announcements from its remote counterpart within a set timeout, it can conclude that the latter is not ready, or temporarily not responsive, and can thus return control to QEGP without wasting entanglement attempts. This MHP synchronization step is also useful for the two sides to verify that they are processing the same QEGP request, and thus catch mismatches.

The MHP synchronization routine inherently incurs some overhead, which is also larger on longer links. We mitigate this overhead by batching entanglement attempts—that is, the physical layer attempts entanglement multiple times after synchronization before reporting back to the link layer. The maximum number of attempts per batch is a purely physical-layer parameter, and it has no relation with the link layer entanglement request timeout parameter described in ref. ¹⁹—although batches should be small enough for the link layer timeout to make sense.

The original design of the QEGP and MHP protocols, as well as our revision, specifies the conceptual interaction between the two protocols and the service exposed to a higher layer in the system, but does not impose particular constraints on how to implement link layer and physical layer, how to realize the physical interface between them, and how to configure things such as the centralized request scheduler and the entanglement attempt procedure. Figure 2 gives an overview of the architecture of our quantum network nodes. We briefly describe our most relevant implementation choices here and in the physical layer section.

Application processing. At the application layer, user programs—written in Python using a dedicated software development kit³⁴—are processed by a rudimentary compilation stage, which translates abstract quantum networking applications into gates and operations supported by our specific quantum physical platform. Such gates and operations are expressed in a low-level assembly-like language for quantum networking applications called NetQASM³⁵. As part of our software stack, we also include an instruction processor, conceptually placed above the link layer, which is in charge of dispatching entanglement requests to QEGP and other application instructions to the physical layer directly.

Interface. Reference ¹⁹ did not provide a specification of the interface to be exposed by the physical layer. We designed this interface such that the physical layer can accept commands from the higher layer, specifically: (a) qubit initialization (INI), (b) qubit measurement (MSR), (c) single-qubit gate (SQG), (d) entanglement attempt (ENT, or ENM for M- or R-type requests), (e) premeasurement gates selection (PMG, to specify in which basis to measure the qubit for M- or R-type requests). For each command, the physical layer reports back an outcome, which indicates whether the command was executed correctly, and can bear the result of a qubit measurement and the Bell state produced after a successful entanglement attempt. Our software stack also comprises a hardware abstraction layer (HAL) that sits below QEGP and the

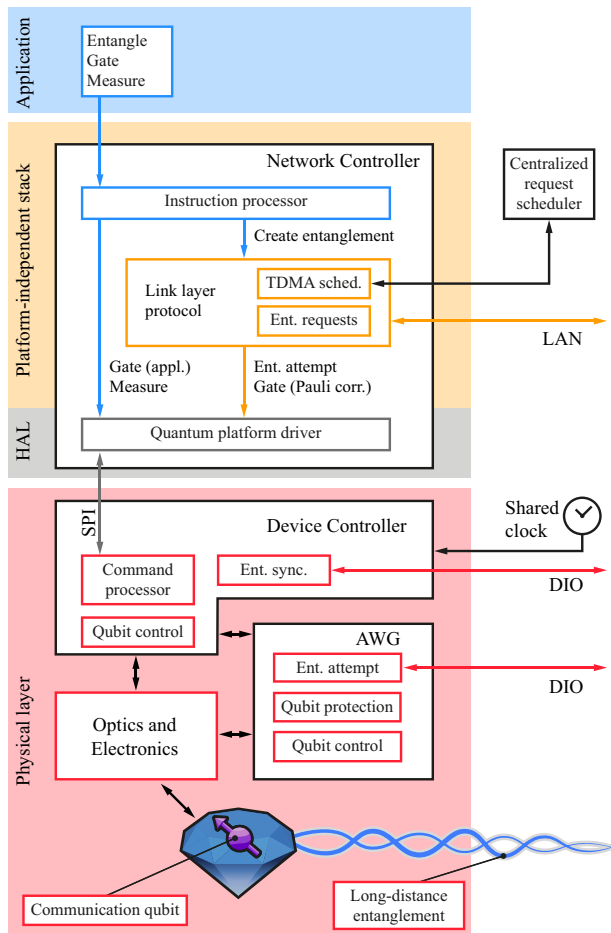


Fig. 2 Quantum network node architecture. From top to bottom: At the application layer, a simple platform-independent routine is sent to the network controller. The network controller implements the platform-independent stack—in this work, only the link layer protocol—and a hardware abstraction layer (HAL) to interface with the physical layer's device controller. An instruction processor dispatches instructions either directly to the physical layer, or to the link layer protocol in case a remote entangled state is requested by the application. The link layer schedules entanglement requests and synchronizes with the remote node (on a local area network, LAN) using a time-division multiple access (TDMA) schedule computed by a centralized scheduler (external). At the physical layer, the device controller fetches commands from—and replies with outcomes to—the network controller. Driven by a clock shared with the neighboring node, it performs hard-real-time synchronization for entanglement generation using a digital input/output (DIO) interface. By controlling the optical and electronic components (among which an arbitrary waveform generator, AWG), the device controller can perform universal quantum control of the communication-qubit in real-time, as well as attempt long-distance entanglement generation with the neighboring node.

instruction processor. The HAL encodes and serializes commands and outcomes, and is thus used to interface with the device controller.

TDMA network schedule. Designing a full-blown centralized request scheduler is a challenge in and of its own, outside the scope of this work. Instead of implementing such a scheduler, we compute static TDMA network schedules²⁶ and install them manually on the two network nodes upon initialization. TDMA schedules for our simple single-link experiments are quite trivial

(see Supplementary Note 1), as the network resources of a node are not contended by multiple links.

Entanglement attempts. Producing entanglement on a link can take several attempts. To minimize the number of ENT commands fetched by MHP from QEGP, as well as to mitigate the MHP synchronization overhead incurred after each entanglement command, we batch entanglement attempts at the MHP layer, such that synchronization and outcome reporting only happens once per batch of attempts.

Delivered entangled states. In our first iteration, we implemented QEGP such that it always delivers $|\Phi^+\rangle$ Bell states to the higher layer. This means that when the physical layer produces a different Bell state, QEGP (on the node where the entanglement request originates) issues a single-qubit gate—a Pauli correction—to transform the entangled pair into the $|\Phi^+\rangle$ state (we abbreviate the four two-qubit maximally entangled Bell states as $|\Phi^\pm\rangle = (|00\rangle \pm |11\rangle)/\sqrt{2}$ and $|\Psi^\pm\rangle = (|01\rangle \pm |10\rangle)/\sqrt{2}$). A future version of QEGP could allow the user to request any Bell state, and could extract the Pauli correction from QEGP so that the application itself can decide, depending on the use case, whether to apply the correction or not.

Mismatch verification. As per our design specification, MHP should also be responsible for verifying that the entanglement commands coming from the two QEGP peers belong to the same request. We did not implement this feature yet because, in our simple quantum network, we do not expect losses on the classical channel used by the two MHP parties to communicate—a lossy classical channel would be the primary source of inconsistencies at the MHP layer¹⁹. However, we believe that this verification step will prove very useful in real-world networks where classical channels do not behave as predictably.

Deployment. We implemented QEGP as a software module in a system that also includes the instruction processor and the hardware abstraction layer. QEGP, the instruction processor and the hardware abstraction layer, forming the network controller, are implemented as a C/C++ standalone runtime developed on top of FreeRTOS, a real-time operating system for embedded platforms³⁶. The runtime and the underlying operating system are deployed on a dedicated Avnet MicroZed—an off-the-shelf platform based on the Zynq-7000 SoC, which hosts two ARM Cortex-A9 processing cores, of which only one is used, clocked at 667 MHz. QEGP connects to its remote peer via TCP over a Gigabit Ethernet interface. The interface to the physical layer is realized through a 12.5 MHz SPI connection. The user application is sent from a general-purpose four-core desktop machine running Linux, which connects to the instruction processor through the same Gigabit Ethernet interface that QEGP uses to communicate with its peer.

Physical layer control in real-time

In this section, we outline the design and operation of the physical layer, which executes the commands issued by the higher layers on the quantum hardware and handles time-critical synchronization between the quantum network nodes. The physical layer of a quantum network, as opposed to the apparatus of a physics experiment, needs to be able to execute commands coming from the layer above in real-time. Additionally, when performing the requested operations, it needs to leave the quantum device in a state that is compatible with future commands (for example, as discussed below, it should protect qubits from decoherence while it awaits further instructions). Finally, if a request cannot be met (e.g., the local quantum hardware is not ready, the remote

quantum hardware is not available, etc.), the physical layer should notify the link layer of the issue without interrupting its service.

Our quantum network is composed of two independent nodes based on diamond NV centers physically separated by ≈ 2 m (see Fig. 2 for the architecture of one node, and Supplementary Fig. 1 for details on the connections between the two nodes). We will refer to the two nodes as client and server, noting that this is only a logical separation useful to describe the case studies—the two nodes have the exact same capabilities. On each node, we implement the logic of the physical layer in a state-machine-based algorithm deployed on a time-deterministic microcontroller, the device controller (Jäger ADwin Pro II, based on Zynq-7000 SoC, dual-core ARM Cortex-A9, clocked at 1 GHz). Additionally, each node uses an arbitrary waveform generator (AWG, Zurich Instruments HDAWG8, 2.4 GSa/s, 300 MHz sequencer) for nanosecond-resolution tasks, such as fast optical and electrical pulses; the use of such a user-programmable FPGA-based AWG, as opposed to a more traditional upload-and-play instrument (such as the ones used in ref. ¹⁵), enables the real-time control of our quantum device.

Single node operation. On our quantum-platform, before a node is available to execute commands, it needs to perform a qubit readiness procedure called charge and resonance check (CR check). This ensures that the qubit system is in the correct charge state and that the necessary lasers are resonant with their respective optical transitions. Other quantum platforms might have a similar preparation step, such as loading and cooling for atoms and ions^{9,10}. Once the CR check is successful, the device controller can fetch commands from the network controller. Depending on the nature of the command, the device controller might need to coordinate with other equipment in the node or synchronize with the device controller of the other node.

For qubit initialization and measurement commands (INI and MSR), the device controller shines the appropriate laser for a pre-defined duration (INI ≈ 100 μ s, MSR ≈ 10 μ s). Both operations are deterministic and carried out entirely by the device controller.

Single-qubit gates (SQG) require the coordination of the device controller and the AWG. For our communication qubits, they consist of generating an electrical pulse with the AWG (duration ≈ 100 ns), which is then multiplied by the qubit frequency (≈ 2 GHz), amplified and finally delivered to the quantum device. The link layer can request rotations in steps of $\pi/16$ around the X, Y or Z axis of the Bloch sphere (here, we implement only X and Y rotations, Z rotations will be implemented in the near future, see Supplementary Note 2). When a new gate is requested by the link layer, the device controller at the physical layer informs the AWG of the gate request via a parallel 32-bit DIO interface. The AWG will then select one of the 64 pre-compiled waveforms, play it, and notify the device controller that the gate has been executed. The device controller will, in turn, notify the network controller of the successful operation.

After the rotation has been performed, our qubit—if left idling—would lose coherence in ≈ 5 μ s. A coherence time exceeding 1 s has been reported on our platform³⁷ using decoupling sequences (periodic rotations of the qubit that shield it from environmental noise). By interleaving decoupling sequences and gates, one can perform extended quantum computations³⁸. These long sequences of pulses have in the past been calculated and optimized offline (on a PC), then uploaded to an AWG, and finally executed on the quantum devices with minimal interaction capabilities (mostly binary branching trees, see ref. ¹⁵). In our case, it is impossible to pre-calculate these sequences, since we cannot know in advance which gates are going to be requested by the link layer. To solve this challenge, we implement a qubit protection module on the AWG, that interleaves decoupling sequences with the requested gates in real-time. As soon as the first gate in a sequence is requested, the AWG starts a decoupling

sequence on the qubit. Then, it periodically checks if a new gate has been requested, and if so, it plays it at the right time in the decoupling sequence. The AWG will continue the qubit protection routine until the device controller will ask for it to stop (e.g., to perform a measurement). This technique allows us to execute universal qubit control without prior knowledge of the sequence to be played and—crucially—in real-time.

Entanglement generation. Differently from the commands previously discussed, attempting entanglement generation (ENT) requires tight timing synchronization between the device controllers—and AWGs—of the two nodes. In our implementation, the two device controllers share a common 1 MHz clock as well as a DIO connection to exchange synchronization messages (see ref. ¹⁵). When the device controllers are booted, they synchronize an internal cycle counter that is used for time-keeping, and is shared, at each node, with their respective network controllers to provide timing information to the link layer and the higher layers. Over larger distances, one could use well-established protocols to achieve sub-nanosecond, synchronized, GPS-disciplined common clocks³⁹.

When a device controller fetches an ENT command, it starts a three-way handshake procedure with the device controller of the other node. If the other node has also fetched an ENT command, they will synchronize and proceed with the entanglement generation procedure. If one of the two nodes is not available (e.g., it is still trying to pass the CR check) the other node will timeout, after 0.5 ms, and return an entanglement synchronization failure (ENT_SYNC_FAIL) to its link layer. The duration of the timeout is chosen such that is comparable with the average time taken by a node to pass the charge and resonance check (if correctly on resonance). This is to avoid unnecessary interactions between the physical layer and the link layer. After the entanglement synchronization step, the device controllers proceed with an optical phase stabilization cycle¹⁵, and then the AWGs are triggered to attempt entanglement generation. In our implementation, one device controller (the server's) triggers both AWGs to achieve sub-nanosecond jitter between the two AWGs (see Supplementary Note 3 for a discussion on longer distance implementation). Each entanglement attempt lasts 3.8 μ s and includes fast qubit initialization, communication-qubit to flying-qubit entanglement, and probabilistic entanglement swapping of the flying qubits¹⁵. The AWGs attempt entanglement up to 1000 times before timing out and reporting an entanglement failure (ENT_FAIL). Longer batches of entanglement attempts would increase the probability that one of the nodes goes into the unwanted charge state (and, therefore, cannot produce entanglement, see Supplementary Note 7). While in principle possible, we did not implement, in this first realization, the charge stabilization mechanism proposed in ref. ¹⁴ that would allow for significantly longer batches of entanglement attempts.

If an entanglement generation attempt is successful (probability $\approx 5 \times 10^{-5}$), the communication qubits of the two nodes will be projected into an entangled state (either $|\Psi^+\rangle$ or $|\Psi^-\rangle$), depending on which detector clicked at the heralding station). To herald the success of the entanglement attempt, a CPLD (Complex Programmable Logic Device, Altera MAX V 5M570ZF256C5N) sends a fast digital signal to both AWGs and device controllers to prevent a new entanglement attempt from being played (which would destroy the generated entangled state). When the heralding signal is detected, the AWGs enter the qubit protection routine and wait for further instructions from the device controllers, which in turn notify the link layer of the successful entanglement generation, as well as which state was generated.

To satisfy M- or R-type entanglement requests, the link layer can instruct the physical layer to apply an immediate measurement to the entangled qubit by means of an ENM command. Up until the

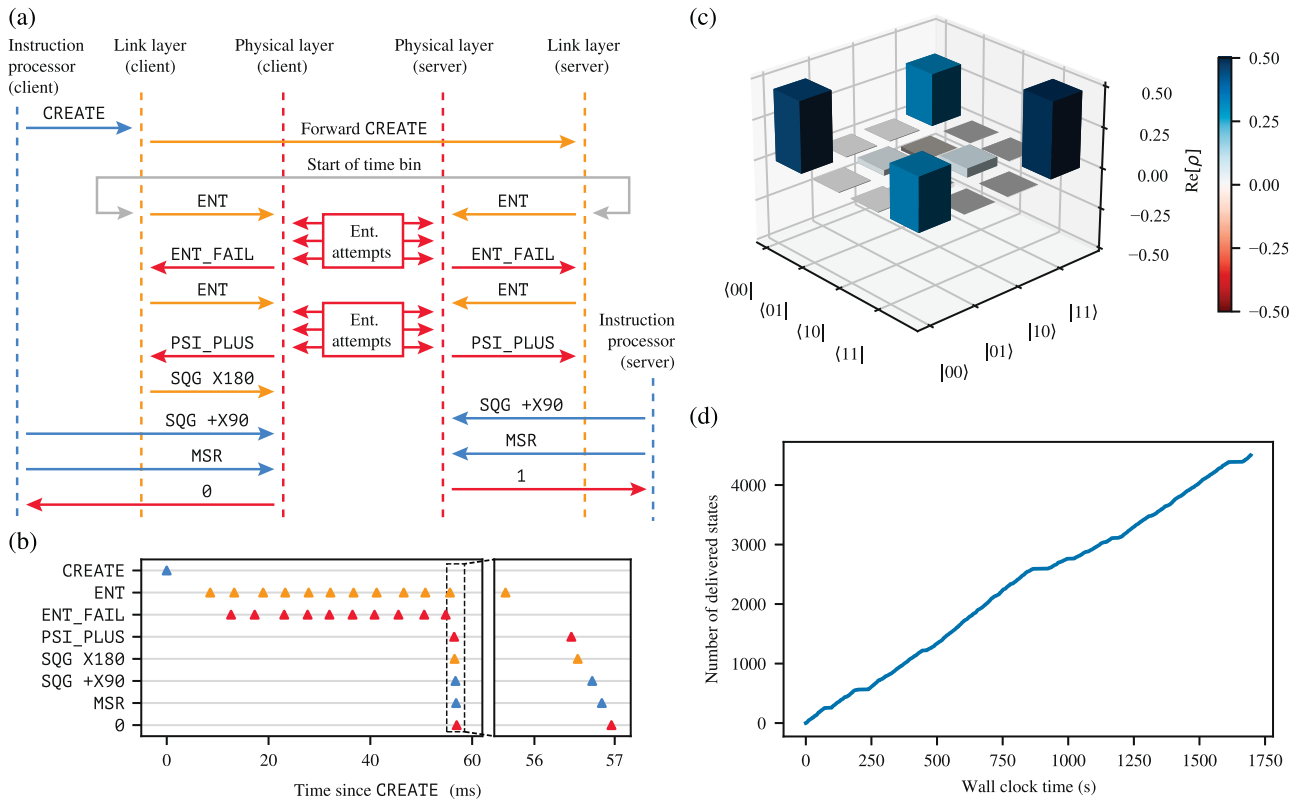


Fig. 3 Full state tomography with the quantum network stack. **a** Sequence diagram of the communication steps across the network stack and the two nodes to perform one repetition of the tomography application (in particular, measurement of the $\langle YY \rangle$ correlator). The coloring follows that of Fig. 1. CREATE: entanglement request, ENT: entanglement attempts request, ENT_FAIL: failed the batch of entanglement attempts, PSI_PLUS: successful entanglement attempt with generated state Ψ^+ , SQG: single-qubit gate, X180: 180° rotation around X axis, MSR: qubit measurement, 0/1: qubit measurement outcome. See Supplementary Table 1 for a complete list of commands. Note that the client's link-layer protocol requests an X180 gate after entanglement generation to deliver the $|\Phi^+\rangle$ Bell state to the higher layer. **b** Example time trace of (a) for the client. Several batches of entanglement attempts are required before an entangled state is heralded. On the right, is a zoomed-in part of the trace (corresponding to the dashed box in the left plot). **c** Reconstructed density matrix of the states delivered by the link layer. Only the real part is plotted (imaginary elements are all ≈ 0 , see main text). We estimate a fidelity F with $|\Phi^+\rangle$ of $F = 0.783(7)$. **d** Total number of delivered states over time. The occasional pauses in entanglement delivery (plateaus) are due to the client's NV center becoming off-resonant with the relevant lasers (see Supplementary Material). Differences in slope are due to changes in resonance conditions that increase the time necessary to pass the charge and resonance check.

heralding of the entangled state, the physical layer operates as it does for the ENT command. When the state is ready, it proceeds immediately with a sequence of single-qubit gates (as prescribed by an earlier PMG command) and a qubit measurement. The result of the measurement, together with which entangled state was generated, is communicated to the link layer. It is worth noting that the two nodes could fetch different types of requests and still generate entanglement. In fact, this will be used later in the remote state preparation application.

Evaluation

To demonstrate and benchmark the capabilities of the link layer protocol, the physical layer, and of our system as a whole, we execute—on our two-node network—three quantum networking applications, all having a similar structure: the client asks for an entangled pair with the server, which QEGP delivers in the $|\Phi^+\rangle$ Bell state, and then both client and server measure their end of the pair in a certain basis. First, we perform full quantum state tomography of the delivered entangled states. Second, we request and characterize entangled states of varying fidelity. Third, we execute remote preparation of qubit states on the server by the client. For all three applications, we study the quality of the entangled pairs delivered by our system.

Additionally, we use the second application to assess the latency incurred by our link layer and to compare it to the overall entanglement generation latency, including that of the physical layer. Crucially, the three applications are executed back-to-back on the quantum network, without any software or hardware changes to the system—the only difference being the quantum-platform-independent application sent to the instruction processor (see Supplementary Note 4).

The sequence diagram in Fig. 3a exemplifies the general flow between system components during the execution of an application. At first, the instruction processor issues a request to create entanglement to the link layer (CREATE). Then, the client's link layer forwards the request to the server's counterpart (Forward CREATE). The request is processed as soon as the designated time bin in the TDMA schedule starts, at which point the first entanglement command (ENT) is fetched by the physical layer. After an entangled state is produced successfully (PSI_PLUS), the link layer of the client issues, if needed, a Pauli correction (π rotation around the X axis, SQG X180) to deliver the pair in the $|\Phi^+\rangle$ state. Finally, the instruction processor issues a gate ($\pi/2$ rotation around the X axis, SQG X90) and a measurement (MSR) to read out the entangled qubit on a certain basis, and receives an outcome from the physical layer (0). Figure 3b illustrates the actual

latencies between these interactions in one iteration of the full state tomography application.

For all our experiments, we configured TDMA time bins to be 20 ms. In a larger network, the duration of time bins should be calibrated according to the average time it takes, on a certain link, to produce an entangled pair of a certain fidelity²⁶. By doing so, one can maximize network usage and thus reduce qubit decoherence on longer end-to-end paths. However, in our single-link network, the duration of time bins only influences the frequency at which new entanglement requests are processed. Our time bin duration accommodates up to four batches of 1000 entanglement attempts.

Full quantum state tomography. The first application consists in generating entangled states at the highest minimum fidelity currently available on our physical setup (0.80), and measuring the two entangled qubits on varying bases to learn their joint quantum state. We measure all 9 two-node correlators ($\langle XX \rangle$, $\langle XY \rangle$, ..., $\langle ZZ \rangle$) as well as all their $\pm$ variations ($\langle +X + X \rangle$, $\langle +X - X \rangle$, etc.) to minimize the bias due to measurement errors. For each of the $9 \times 4 = 36$ combinations, we measure 125 data points, for a total of 4500 entangled states generated and measured. The Supplementary Material contains a pseudocode description of the application.

The collected measurement outcomes are then analyzed using QInfer⁴⁰, in particular, the Monte Carlo method described in ref. ⁴¹ for Bayesian estimation of density matrices from tomographic measurements. The reconstructed density matrix is displayed in Fig. 3c (only the real part is shown in the figure) and its values and uncertainties (1 s.d.) are

$$\text{Re}[\rho] = \begin{pmatrix} 0.442(6) & 0.003(3) & 0.003(2) & 0.328(5) \\ 0.003(3) & 0.033(6) & -0.023(5) & -0.000(5) \\ 0.003(2) & -0.023(5) & 0.056(4) & -0.003(4) \\ 0.328(5) & -0.000(5) & -0.003(4) & 0.469(7) \end{pmatrix},$$

$$\text{Im}[\rho] = \begin{pmatrix} 0 & -0.014(3) & -0.005(7) & 0.032(5) \\ 0.014(3) & 0 & -0.002(4) & 0.001(5) \\ 0.005(7) & 0.002(4) & 0 & -0.000(7) \\ -0.032(5) & -0.001(5) & 0.000(7) & 0 \end{pmatrix}.$$

Here $\rho_{ij,mn} = \langle ij | \rho | mn \rangle$, with i, m (j, n) being the client (server) qubit states in the computational basis. The uncertainty on each element of the density matrix is calculated as the standard deviation of that element over the probability distribution approximated by the Monte Carlo reconstruction algorithm (probability distribution approximated by 10^5 Monte Carlo particles⁴¹). It is then possible to estimate the fidelity of the delivered entangled states with respect to the maximally entangled Bell state, which we find to be $F = 0.783(7)$. The measured fidelity is slightly lower ($\approx 3\%$) than what was measured in ref. ¹⁵ without the use of the QEGP abstraction (and the whole network controller where QEGP runs). This discrepancy could be due to the additional physical-layer decoupling sequences required for real-time operation ($\approx 300 \mu\text{s}$) and the additional single-qubit gate issued by the link layer to always deliver $|\Phi^+\rangle$ (see Supplementary Note 5).

It is to be noted that, in order to obtain the most faithful estimate of the generated state (see Supplementary Note 6 for details), the measured expectation values are corrected, in post-processing, to remove known tomography errors of both client and server⁴², and events in which at least one physical device was in the incorrect charge state.

Finally, we show, in Fig. 3d, that our system can sustain a fairly stable entanglement delivery rate over ≈ 30 min of data acquisition—plateaus and changes in slope can be attributed to varying conditions of resonance between the NV centers and the relevant lasers (see Supplementary Material).

Latency vs fidelity. The QEGP interface allows its user to request entangled pairs at various minimum fidelities. For physical reasons, higher fidelities will result in lower entanglement generation rates^{14,17}. The trade-off between fidelity and throughput is particularly interesting in a scenario where some applications might require high-fidelity entangled pairs and are willing to wait a long time, while others might prefer lower-fidelity states but higher rates¹⁹. Clearly, for the link layer to offer a range of fidelities to choose from, the underlying physical layer must support such a range. We benchmark the capabilities of the link layer and of the physical layer to deliver states at various fidelities in a single application by measuring the $\langle XX \rangle$, $\langle YY \rangle$ and $\langle ZZ \rangle$ correlators (and their $\pm$ variations, as we did above, for a total of $3 \times 4 = 12$ correlators) for seven different target fidelities, (0.50, 0.55, 0.60, 0.65, 0.70, 0.75, 0.80). We generate 1500 entangled states per fidelity, for a total of 10,500 delivered states (see Supplementary Material). With this case study, we analyze both the resulting fidelity and the system's latency for different requested fidelities.

The results for measured fidelity versus requested fidelity are shown in Fig. 4a. It is worth noting that the application iterates over the range of fidelities in real-time, and thus the physical layer is prepared to deliver any of them at any point. We calibrate the physical layer to deliver states of slightly higher fidelity than the requested ones (0.03 more), since entanglement requests specify the minimum desired fidelity. The measured fidelities are—within measurement uncertainty—always matching or exceeding the requested minimum ones (the dashed gray line in Fig. 4a is the $y=x$ diagonal). As in the previous application, measurement outcomes are post-processed to eliminate tomography errors and events in which the physical devices were in the incorrect charge state (we refer to the latter as charge state correction). For arbitrary applications that use the delivered entangled states for something other than statistical measurements, applying the second correction directly at the link layer might prove challenging, since the information concerning whether to discard an entangled pair is only available at the physical layer after the entangled state is delivered to the link layer (when the next CR check is performed). However, a mechanism to identify badly entangled pairs retroactively at the link layer—like the expiry functionality included in the original design of QEGP¹⁹—could be used to discard entangled states after they have been delivered by the physical layer. For completeness, we also report, again in Fig. 4a, the measured fidelity when the wrong charge state correction is not applied.

For each requested fidelity, we also measure the entanglement generation latency¹⁹, defined as the time between the issuing of the CREATE request to the link layer, until the successful entanglement outcome reported by the physical layer (refer to Fig. 3a for a diagram of the events in between these two). Figure 4b shows the measured average latency, grouped by requested fidelity and broken down into various sources of latency. When calculating the average latencies, we have ignored entanglement requests that required more than 10 s to be fulfilled. These high-latency requests correspond to the horizontal plateaus of Fig. 3d (see Supplementary Material for details). The main contribution to the total latency comes from the entanglement generation process at the physical layer, followed by the NV center preparation time (CR check). Both latency values are consistent with the expected number of entanglement attempts required by the single-photon entanglement protocol employed at the physical layer¹⁴. The link layer protocol adds, on average, ≈ 10 ms of extra latency to all requests, regardless of their fidelity. This is due partly to the synchronization of the CREATE request between the two nodes (i.e. a simple TCP message), but mostly to the nodes having to wait for the next time bin in the network schedule to start (the larger the time bins, the larger the worst-case waiting time, see Supplementary Material). We remark that,

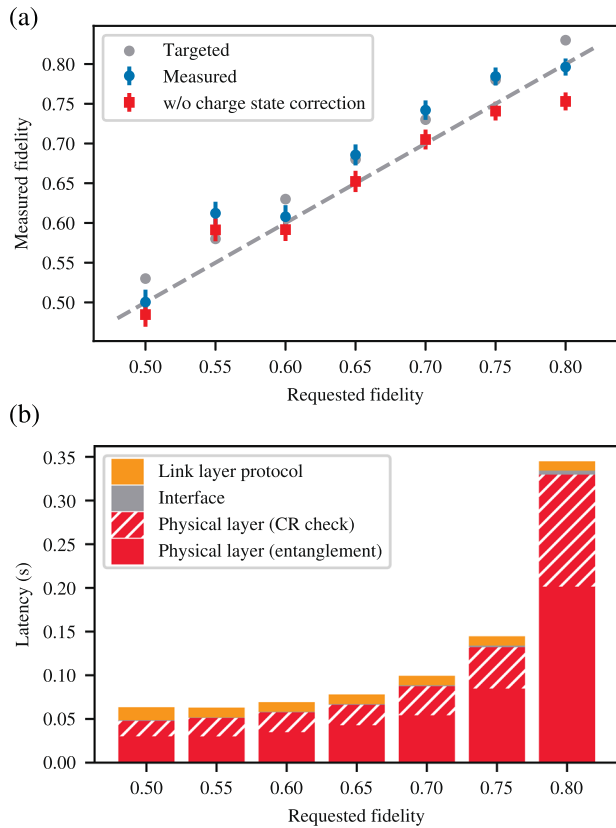


Fig. 4 Performance of the entanglement delivery service. **a** Measured fidelity of the states delivered by the link layer for varying requested fidelity. Targeted fidelity at the physical layer is 0.03 higher than the link layer protocol's minimum fidelity request. When not correcting for wrong charge state events, fidelity is reduced by a few percents (see Supplementary Material). Error bars represent 1 s.d. **b** Average latency of the entanglement delivery per requested fidelity, broken down into sources of latency. Entanglement generation and charge and resonance check at the physical layer are the largest sources of latency (at higher fidelities, more entanglement attempts are required before success). Running the link layer protocol introduces a small but measurable overhead (≈ 10 ms) to the entanglement generation procedure, which does not depend on the requested fidelity, and that could be mitigated by requesting multiple entangled states in a single instruction. The communication delays between the quantum network controller and quantum device controller (Interface) introduce negligible overall latency.

by requesting multiple entangled states in a single CREATE, one can distribute this overhead over many generated pairs, to the point where it becomes negligible. While our applications did not issue multi-pair CREATE requests, this would be the more natural choice for real applications and would result in better utilization of the allocated time bins. Finally, the overhead incurred by the interface between microcontrollers is rather small (barely visible in Fig. 4b) but could, however, be further reduced by integrating the device controller and network controller into a single device. It is worth mentioning that, in our simple scenario in which each entanglement request is only submitted to QEGP after the previous one completes, and thus the request queue never grows larger than one element, throughput happens to be almost exactly the same as the inverse of latency, and hence it is not reported here.

Overall, we observe that the extra entanglement generation latency incurred when deploying an abstraction layer (QEGP) on top of the physical layer, while not too modest, is only a small part of the whole, particularly at higher fidelities. Nevertheless,

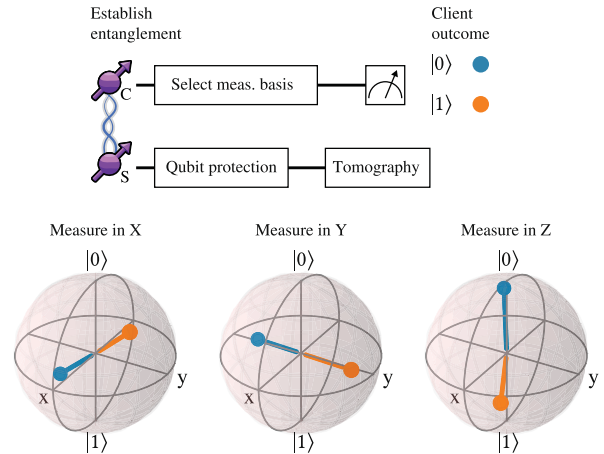


Fig. 5 Tomography of states prepared on the server by the remote client. For each chosen measurement axis of the client (X, Y, Z), and for each obtained measurement outcome at the client ($|0\rangle$, $|1\rangle$), a different state is prepared on the server. Plotted on the Bloch spheres are the results of the tomography on the server's qubit. Uncertainties on each coordinate are ≈ 0.05 (see Supplementary Material). We find an average preparation fidelity of $F = 0.853(8)$.

optimizing the length of TDMA time bins could result in an even smaller overhead (see Supplementary Material).

Remote state preparation. One of the use cases of the QEGP service is to prepare quantum states on a remote quantum server¹⁹. Remote state preparation is a fundamental step to execute a blind quantum computation application⁵, whereby a client quantum computer with limited resources can run quantum applications on a powerful remote quantum server using the many qubits the server has, while keeping the performed computation private.

Remote state preparation is different from the previous two cases in that the client can measure its end of the entangled pair as soon as the pair is generated, while the server has to keep its qubit alive, waiting for further instructions. For such a scenario, the client can make use of QEGP's service to issue R-type entanglement requests, so that the local end of the entangled pair can be measured (on a certain basis) as soon as it is generated, while the server's qubit can be protected for later usage. An R-type entanglement request results in an ENM command on the client and an ENT command on the server. For this type of request (as well as for M-type ones), since the local end of the pair is measured immediately, the client's QEGP can skip the Pauli correction used to always deliver $|\Phi^+\rangle$, and can instead apply a classical correction to the received measurement outcome (see Supplementary Material).

To showcase this feature of QEGP, we use the client node to prepare the six cardinal states on the server ($|\pm x\rangle$, $|\pm y\rangle$, $|0\rangle$ and $|1\rangle$) by having the client measure its share of the entangled state in the six cardinal bases. We then let the server measure the prepared states—again in the six cardinal bases—to perform tomography. For each client measurement basis and for each server tomography basis, we deliver 125 entangled states at a requested fidelity of 0.80, for a total of $6 \times 6 \times 125 = 4500$ remote state preparations (see Supplementary Material). The results are presented in Fig. 5, which displays the tomography of the prepared states on the server for the three different measurement axes of the client and the two possible measurement outcomes of the client. The prepared states are affected by the measurement error of the client ($F_0 = 0.928(3)$, $F_1 = 0.997(1)$): an error in the measurement of the client's qubit results in an incorrect identification of the state prepared on the server. By alternating

between positive and negative readout orientations, we make sure that the errors affect all prepared states equally instead of biasing the result. We note that we exclude, once again, events in which at least one of the two devices was in the wrong charge state, and we correct for the known tomography error on the server (results without corrections are in the Supplementary Material). Overall, we find an average remote state preparation fidelity of $F = 0.853(8)$. The asymmetry in the fidelity of the $|0\rangle$ and $|1\rangle$ states is caused by the asymmetry in the populations $\langle 01|\rho|01\rangle$ vs $\langle 10|\rho|10\rangle$ of the delivered entangled state, which in turn is due to the double $|0\rangle$ occupancy error of the single-photon protocol used to generate entanglement^{14,15}.

DISCUSSION

In summary, we have demonstrated the operation of a link layer and a physical layer for entanglement-based quantum networks. The link layer abstracts the entanglement generation procedure provided by the physical layer—implemented here with two NV center-based quantum network nodes—into a robust platform-independent service that can be used to run quantum networking applications. We performed full quantum state tomography of the states delivered by the link layer, tested its ability to deliver states at different fidelities in real-time, and verified remote state preparation of a qubit from the client on the server, a fundamental step towards blind quantum computation⁵. We have shown that our implementation of link and physical layers can deliver entangled states at the fidelity requested by the user, despite some marginal inefficiencies—some of which can be addressed in a future version of the protocols (e.g., avoiding Pauli corrections unless necessary). We have also quantified the additional latency incurred by deploying the link layer protocol on top of the physical layer. Although not detrimental, the extra overhead is still noticeable, but can also be scaled down by optimizing the scheduling of entanglement generation requests. We also acknowledge that scheduling a quantum node's resources is still an open problem^{26,43,44} and that the simple approach taken here is likely a suboptimal choice for more advanced quantum networks. We emphasize, however, that our link layer protocol is not tied to any particular scheduling algorithm or architecture—it merely expects that the schedule of each node be matched with its peer. In ref. ¹⁹ for example, the schedule was instead formed via a distributed queue protocol, and in the future other architectures and algorithms²⁶ may be more suitable for scaling to larger networks.

Other research challenges posed by our work include an in-depth analysis of the security of quantum network implementations. For example, it is clear that if the classical control messages used in our protocol are not authenticated, unwanted entanglement generation may be triggered at one of the nodes. In some physical layer implementations such as the one considered here, this may negatively impact the quality of the qubits already stored at the node⁴⁵, and hence impact availability. Initial work indicates that the performance impact of adding authentication, however, is small (less than a 3% reduction in throughput in a non-optimized simulation study)⁴⁶.

The adoption of the techniques presented here (which are not specific to our diamond devices) by other quantum network platforms^{9,10,17,47–50} will boost the development of large-scale and heterogeneous quantum networks. Real-time control of memory qubits, as well as the availability of multi-node networks and dynamic network schedules, will enable demonstrations of the higher layers of the network stack⁵¹, which in turn will open the door to end-to-end connectivity on a platform-independent quantum network.

DATA AVAILABILITY

The datasets that support this manuscript and the software to analyze them are available at 4TU.ResearchData with the identifier <https://doi.org/10.4121/16912522>.

CODE AVAILABILITY

The application software development kit is open-sourced on GitHub³⁴. The link layer implementation is part of a larger software project to develop an operating system for quantum network nodes that is still under development and thus not currently available for external use.

Received: 25 November 2021; Accepted: 15 September 2022;

Published online: 15 October 2022

REFERENCES

- Kimble, H. J. The quantum internet. *Nature* **453**, 1023 (2008).
- Wehner, S., Elkouss, D. & Hanson, R. Quantum internet: a vision for the road ahead. *Science* **362**, eaam9288 (2018).
- Ekert, A. & Renner, R. The ultimate physical limits of privacy. *Nature* **507**, 443 (2014).
- Jiang, L., Taylor, J. M., Sørensen, A. S. & Lukin, M. D. Distributed quantum computation based on small quantum registers. *Phys. Rev. A* **76**, 062323 (2007).
- Broadbent, A., Fitzsimons, J. & Kashefi, E. Universal blind quantum computation. In *Proc. 50th Annual IEEE Symposium on Foundations of Computer Science* 517–526 (2009).
- Gottesman, D., Jennewein, T. & Croke, S. Longer-baseline telescopes using quantum repeaters. *Phys. Rev. Lett.* **109**, 070503 (2012).
- Kómár, P. et al. A quantum network of clocks. *Nat. Phys.* **10**, 582 (2014).
- Moehring, D. L. et al. Entanglement of single-atom quantum bits at a distance. *Nature* **449**, 68 (2007).
- Stephenson, L. et al. High-rate, high-fidelity entanglement of qubits across an elementary quantum network. *Phys. Rev. Lett.* **124**, 110501 (2020).
- Ritter, S. et al. An elementary quantum network of single atoms in optical cavities. *Nature* **484**, 195 (2012).
- Hofmann, J. et al. Heralded entanglement between widely separated atoms. *Science* **337**, 72 (2012).
- Bernien, H. et al. Heralded entanglement between solid-state qubits separated by three metres. *Nature* **497**, 86 (2013).
- Kalb, N. et al. Entanglement distillation between solid-state quantum network nodes. *Science* **356**, 928 (2017).
- Humphreys, P. C. et al. Deterministic delivery of remote entanglement on a quantum network. *Nature* **558**, 268 (2018).
- Pompili, M. et al. Realization of a multinode quantum network of remote solid-state qubits. *Science* **372**, 259 (2021).
- Delteil, A. et al. Generation of heralded entanglement between distant hole spins. *Nat. Phys.* **12**, 218 (2016).
- Stockill, R. et al. Phase-tuned entangled state generation between distant spin qubits. *Phys. Rev. Lett.* **119**, 010503 (2017).
- Aparicio, L., Van Meter, R. & Esaki, H. Protocol design for quantum repeater networks. In *Proc. 7th Asian Internet Engineering Conference AINTEC '11* 73–80 (Association for Computing Machinery, 2011).
- Dahlberg, A. et al. A link layer protocol for quantum networks. In *Proc. ACM Special Interest Group on Data Communication, SIGCOMM '19* 159–173 (Association for Computing Machinery, 2019).
- Pirker, A. & Dür, W. A quantum network stack and protocols for reliable entanglement-based networks. *N. J. Phys.* **21**, 033003 (2019).
- Kozłowski, W., Dahlberg, A. & Wehner, S. Designing a quantum network protocol. In *Proc. 16th International Conference on Emerging Networking Experiments and Technologies, CoNEXT '20* 1–16 (Association for Computing Machinery, 2020).
- Aguado, A. et al. Enabling quantum key distribution networks via software-defined networking. In *International Conference on Optical Network Design and Modeling (ONDM)* 1–5 (Association for Computing Machinery, 2020).
- Kozłowski, W., Kuipers, F. & Wehner, S. A p4 data plane for the quantum internet. In *Proc. 3rd P4 Workshop in Europe, EuroP4'20* 49–51 (Association for Computing Machinery, 2020).
- Alshowkan, M. et al. Reconfigurable quantum local area network over deployed fiber. *PRX Quantum* **2**, 040304 (2021).
- Ferguson, A. D. et al. Orion: Google's software-defined networking control plane. In *18th USENIX Symposium on Networked Systems Design and Implementation (NSDI 21)* 83–98 (2021).
- Skrzypczyk, M. & Wehner, S. An architecture for meeting quality-of-service requirements in multi-user quantum networks. Preprint at <http://arxiv.org/abs/2111.13124> (2021).

27. Van Meter, R., Satoh, T., Ladd, T. D., Munro, W. J. & Nemoto, K. Path selection for quantum repeater networks. *Netw. Sci.* **3**, 82 (2013).
28. Caleffi, M. Optimal routing for quantum networks. *IEEE Access* **5**, 22299 (2017).
29. Gyongyosi, L. & Imre, S. Decentralized base-graph routing for the quantum internet. *Phys. Rev. A* **98**, 022310 (2018).
30. Pant, M. et al. Routing entanglement in the quantum internet. *npj Quantum Inf.* **5**, 1 (2019).
31. Chakraborty, K., Rozpedek, F., Dahlberg, A. & Wehner, S. Distributed routing in a quantum internet. Preprint at <http://arxiv.org/abs/1907.11630> (2019).
32. Shi, S. & Qian, C. Concurrent entanglement routing for quantum networks: model and designs. In *Proc. Annual Conference of the ACM Special Interest Group on Data Communication on the Applications, Technologies, Architectures, and Protocols for Computer Communication*, SIGCOMM '20 62–75 (Association for Computing Machinery, 2020).
33. Chakraborty, K., Elkouss, D., Rijsman, B. & Wehner, S. Entanglement distribution in a quantum network: a multicommodity flow-based approach. *IEEE Trans. Quantum Eng.* **1**, 1 (2020).
34. NetQASM SDK. <https://github.com/QuTech-Delft/netqasm>.
35. Dahlberg, A. et al. NetQASM-a low-level instruction set architecture for hybrid quantum-classical programs in a quantum internet. *Quantum Sci. Technol.* **7**, 035023 (2022).
36. FreeRTOS. Real-time operating system for microcontrollers. <https://www.freertos.org/> (2021).
37. Abobeih, M. H. et al. One-second coherence for a single electron spin coupled to a multi-qubit nuclear-spin environment. *Nat. Commun.* **9**, 2552 (2018).
38. Bradley, C. et al. A ten-qubit solid-state spin register with quantum memory up to one minute. *Phys. Rev. X* **9**, 031045 (2019).
39. Serrano, J. et al. The white rabbit project. <https://white-rabbit.web.cern.ch/> (2013).
40. Granade, C. et al. QInfer: statistical inference software for quantum applications. *Quantum* **1**, 5 (2017).
41. Granade, C., Combes, J. & Cory, D. G. Practical Bayesian tomography. *N. J. Phys.* **18**, 033024 (2016).
42. Nachman, B., Urbanek, M., de Jong, W. A. & Bauer, C. W. Unfolding quantum computer readout noise. *npj Quantum Inf.* **6**, 1 (2020).
43. Vardoyan, G., Guha, S., Nain, P. & Towsley, D. On the stochastic analysis of a quantum entanglement distribution switch. *IEEE Trans. Quantum Eng.* **2**, 1 (2021).
44. Vardoyan, G., Guha, S., Nain, P. & Towsley, D. On the capacity region of bipartite and tripartite entanglement switching. *SIGMETRICS Perform. Eval. Rev.* **48**, 45–50 (2021).
45. Kalb, N., Humphreys, P. C., Slim, J. J. & Hanson, R. Dephasing mechanisms of diamond-based nuclear-spin memories for quantum networks. *Phys. Rev. A* **97**, 062330 (2018).
46. Abrahams, J. *Data Origin Authentication and the Classical Data Plane of a Quantum Network Link*. Master's thesis, Delft Univ. Technology (2022).
47. Rose, B. C. et al. Observation of an environmentally insensitive solid-state spin defect in diamond. *Science* **361**, 60 (2018).
48. Nguyen, C. et al. Quantum network nodes based on diamond qubits with an efficient nanophotonic interface. *Phys. Rev. Lett.* **123**, 183602 (2019).
49. Trusheim, M. E. et al. Transform-limited photons from a coherent tin-vacancy spin in diamond. *Phys. Rev. Lett.* **124**, 023602 (2020).
50. Son, N. T. et al. Developing silicon carbide for quantum spintronics. *Appl. Phys. Lett.* **116**, 190501 (2020).
51. Kozłowski, W. & Wehner, S. Towards large-scale quantum networks. In *Proc. Sixth Annual ACM International Conference on Nanoscale Computing and Communication*, NANOCOM '19 1–7 (Association for Computing Machinery, 2019).

ACKNOWLEDGEMENTS

We thank Joris van Rantwijk, Sidney Cadot, Ludo Visser, and Nicolas Demetriadou for experimental support, Nico Seidler and Önder Karpas for useful discussions, and Simon Baier for comments on an earlier version of this manuscript. We acknowledge financial support from the EU Flagship on Quantum Technologies through the project Quantum Internet Alliance (EU Horizon 2020, grant agreement no. 820445); from the Netherlands Organization for Scientific Research (NWO) through the Zwaartekracht program Quantum Software Consortium (project no. 024.003.037/3368); from the European Research Council (ERC) through a Consolidator Grant (grant agreement no. 772627 to R.H.) under the European Union's Horizon 2020 Research and Innovation Program.

AUTHOR CONTRIBUTIONS

M.P., C.D.D., I.t.R., W.K., R.H. and S.W. devised the experiment. M.P., C.D.D., I.t.R. and L.d.K. carried out the experiments and collected the data. M.P., L.d.K., A.J.S. and S.L.N.H. prepared the quantum-platform apparatus. C.D.D., I.t.R., B.v.d.V., M.S., G.F., P.P. and W.K. prepared the platform-independent apparatus. M.P., C.D.D., R.H. and S.W. wrote the manuscript and the supplementary materials with input from all authors. M.P., C.D.D., I.t.R. and L.d.K. analyzed the data and discussed them with all authors. M.P. and C.D.D. contributed equally to this work. R.H. and S.W. supervised the research.

COMPETING INTERESTS

The authors declare no competing interests.

ADDITIONAL INFORMATION

Supplementary information The online version contains supplementary material available at <https://doi.org/10.1038/s41534-022-00631-2>.

Correspondence and requests for materials should be addressed to R. Hanson or S. Wehner.

Reprints and permission information is available at <http://www.nature.com/reprints>

Publisher's note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.



Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this license, visit <http://creativecommons.org/licenses/by/4.0/>.

© The Author(s) 2022