

A crossbar network for silicon quantum dot qubits

Li, Ruoyu; Petit, Luca; Franke, David P.; Dehollain, Juan Pablo; Helsen, Jonas; Steudtner, Mark; Thomas, Nicole K.; Wehner, Stephanie; Vandersypen, Lieven M.K.; Veldhorst, Menno

DOI

[10.1126/sciadv.aar3960](https://doi.org/10.1126/sciadv.aar3960)

Publication date

2018

Document Version

Final published version

Published in

Science Advances

Citation (APA)

Li, R., Petit, L., Franke, D. P., Dehollain, J. P., Helsen, J., Steudtner, M., Thomas, N. K., Wehner, S., Vandersypen, L. M. K., & Veldhorst, M. (2018). A crossbar network for silicon quantum dot qubits. *Science Advances*, 4(7), Article eaar3960. <https://doi.org/10.1126/sciadv.aar3960>

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.

CONDENSED MATTER PHYSICS

A crossbar network for silicon quantum dot qubits

Ruoyu Li^{1,2}, Luca Petit^{1,2}, David P. Franke^{1,2}, Juan Pablo Dehollain^{1,2}, Jonas Helsen¹, Mark Steudtner^{3,1}, Nicole K. Thomas⁴, Zachary R. Yoscovits⁴, Kanwal J. Singh⁴, Stephanie Wehner¹, Lieven M. K. Vandersypen^{1,2,4}, James S. Clarke⁴, Menno Veldhorst^{1,2*}

The spin states of single electrons in gate-defined quantum dots satisfy crucial requirements for a practical quantum computer. These include extremely long coherence times, high-fidelity quantum operation, and the ability to shuttle electrons as a mechanism for on-chip flying qubits. To increase the number of qubits to the thousands or millions of qubits needed for practical quantum information, we present an architecture based on shared control and a scalable number of lines. Crucially, the control lines define the qubit grid, such that no local components are required. Our design enables qubit coupling beyond nearest neighbors, providing prospects for nonplanar quantum error correction protocols. Fabrication is based on a three-layer design to define qubit and tunnel barrier gates. We show that a double stripline on top of the structure can drive high-fidelity single-qubit rotations. Self-aligned inhomogeneous magnetic fields induced by direct currents through superconducting gates enable qubit addressability and readout. Qubit coupling is based on the exchange interaction, and we show that parallel two-qubit gates can be performed at the detuning-noise insensitive point. While the architecture requires a high level of uniformity in the materials and critical dimensions to enable shared control, it stands out for its simplicity and provides prospects for large-scale quantum computation in the near future.

INTRODUCTION

The widespread interest in quantum computing has motivated the development of conceptual architectures across a range of disciplines (1–5). Remarkable differences between the various approaches become apparent when considering the physical size of the qubit. A recent proposal for a microwave-trapped ion quantum computer hosting 2 billion qubits puts the required area to an astonishing size of more than 100 m × 100 m (5). The same number of superconducting qubits is estimated to require an area of 5 m × 5 m (2). Qubits defined by the spin states of semiconductor quantum dots, on the other hand, could fit in an area of less than 5 mm × 5 mm. Small components can provide essential benefits in terms of scalability, and effort to demonstrate the physical operation has already culminated in the realization of high-fidelity single-qubit rotations, two-qubit logic gates, small quantum algorithms, and simple quantum error correction schemes (6–10). However, as current qubit technology requires control lines for every qubit, a key challenge is to avoid an interconnect bottleneck for full control over a large qubit grid (11).

Conventional processors can have more than 2 billion transistors on a 21.5 × 32.5-mm² die (12). Such a high packaging density crucially relies on a limited number of input-output connections (IOs). Transistor-to-IO ratios can be as high as 10⁶ (11) because of an integration of the so-called crossbar technology. Combinations of row lines (RLs) and column lines (CLs) enable the identification of unique points on a grid structure, providing a mechanism for large-scale parallel and rapid read/write instructions. In decades of advancements in semiconductor technology, this concept has resulted in today's most powerful supercomputers. Recently, the idea to implement similar shared control schemes for quantum systems has been introduced for donor-based systems (3) and later proposed for quantum dot spins (4, 11). In the work by Hill *et al.* (3), qubits are defined on the nuclear

spin states of phosphorus donors in silicon, and a scheme is introduced where electrons can be shuttled to and from the donor using shared control. The change in charge occupation after shuttling shifts the resonance frequency of the nuclear spin qubit and thus provides qubit addressability. In the work by Veldhorst *et al.* (4), floating gates addressed via transistor circuits connected to a crossbar array control quantum dot qubits. This stimulated early proof-of-principle operations, such as local transistor-controlled charge detection (13), but requires extensive developments in downscaling and developing new devices such as vertical transistors. Thus, while both proposals offer the prospect of a significant reduction in the number of connections to external control logic, they also rely on feature sizes and integration schemes that are not compatible with today's industry standards and that are far beyond current experimental capabilities.

Here, we propose a crossbar scheme for a two-dimensional (2D) array of quantum dots that can operate a large number of qubits with high fidelity. The structure is simple and elegant in design and does not require extremely small feature sizes, but instead relies on a high level of uniformity. Specifically, we require that a single voltage applied to a common gate can bring individual dots to the single-electron occupancy. In addition, depending on the operation mode, we require that the variation of tunnel coupling between quantum dots can be engineered to be within one order of magnitude. Continuous progress in fabrication has already led to individual double-dot systems with this level of charge uniformity (14, 15). We envisage that metrics, such as variation in threshold voltage, charging energy, and tunnel coupling, will need to improve by approximately an order of magnitude to use common gates in large quantum dot grids, and a promising platform to achieve this is advanced semiconductor manufacturing. Building upon these arrays, we introduce a spin qubit module that combines global charge control, local tunability, and electron shuttling between dots with alternating local magnetic fields and global electron spin resonance (ESR) control. Truly large-scale quantum computing can be achieved by connecting multiple of these qubit modules. We will conclude by providing an overview of the challenges and opportunities for quantum algorithms and quantum error correction on the crossbar spin qubit network.

Copyright © 2018
The Authors, some
rights reserved;
exclusive licensee
American Association
for the Advancement
of Science. No claim to
original U.S. Government
Works. Distributed
under a Creative
Commons Attribution
NonCommercial
License 4.0 (CC BY-NC).

¹QuTech, Delft University of Technology, Lorentzweg 1, 2628 CJ Delft, Netherlands.

²Kavli Institute of Nanoscience, University of Technology, P.O. Box 5046, 2600 GA Delft, Netherlands. ³Instituut-Lorentz, Universiteit Leiden, P.O. Box 9506, 2300 RA Leiden, Netherlands. ⁴Components Research, Intel Corporation, 2501 Northwest 229th Avenue, Hillsboro, OR 97124, USA.

*Corresponding author. Email: m.veldhorst@tudelft.nl

RESULTS

Crossbar network layout

Figure 1 schematically shows the gate layout of the qubit module containing a 2D quantum dot array. The qubits are based on the spin states of single electrons that are induced by electric gates in isotopically purified silicon (^{28}Si) quantum dots, reducing decoherence due to nuclear spin noise (7). The architecture is agnostic to the integration scheme, and the quantum dots can be located at a Si/SiO₂ interface (7), where the abrupt change in band structure can cause a large valley splitting energy, leading to well-isolated qubit states. Alternatively, the quantum dots can be formed in a Si/SiGe quantum well stack (6), where the epitaxial nature of the SiGe interface may be beneficial to meet the required uniformity for global operation as considered in the architecture here.

The architecture consists of a crossbar gate structure of three in-plane layers (see Fig. 1, A and B, and superconducting striplines on top). The striplines deliver global radio frequency (RF) pulses to manipulate the spin state, as will be discussed below. The first layer hosts the CLs, which supply voltages to the horizontal barrier gates. The CLs also carry direct currents (DC) for the generation of the magnetic field pattern (see also Fig. 2, C and D). These gates are deposited as the first layer to accommodate a well-defined cross section and are made of superconducting material. The subsequent RLs are isolated from the first layer of gates and supply the voltages to the vertical barrier gates. The plunger gates are formed through vias that connect to the qubit lines (QLs). This gating scheme does not require smaller manufacturing elements than the quantum dots and the interdot tunnel barriers. Here, we consider barrier and plunger gate widths of 30 and 40 nm, respectively, and quantum dot pitch spacing of 100 nm. These numbers enable more than 1000 qubits to fit in an area smaller than $5\ \mu\text{m} \times 5\ \mu\text{m}$ (note that, in our architecture, half of the quantum dots host a qubit, increasing the area by a factor of 2). These dimensions are compatible with advanced semiconductor manufacturing and multiple patterning (16).

Figure 1B shows a conceptual image of a qubit module. In the idle state, each qubit has four empty neighboring dots. This is achieved by setting the bias voltages applied to the diagonal qubit gates, alternating between accumulation and depletion modes. This sparse occupation has several advantages: It increases the number of control gates per qubit without changing the physical gating density, the sparsely spaced qubits reduce crosstalk, and the empty sites will

enable the shuttling of qubits between different sites. The gate pattern allows for selective addressing of qubits with the combined operation of the different gate layers, as discussed below. For N qubits occupying a square dot array, the combined control reduces the total number of gate lines to $N_{\text{totalwires}} \approx 4\sqrt{(2N)} + 1$. The analog control signals can be fed through the qubit network at the periphery, and no additional control elements are needed within the grid. This allows for a dense packing of the quantum dots.

Since each gate is shared by a line of quantum dots, a high level of uniformity across the whole structure is required. These requirements can, however, be relaxed significantly when aiming for a parallel qubit operation in a line-by-line manner. Here, the long coherence times of silicon qubits become crucial (7). We require that the tunnel coupling t_0 be globally controlled to below 10 Hz in the off state and in the range of 10 to 100 GHz in the on state, depending on the operation mode. The lower bound is set by the error threshold due to unwanted shuttling during a quantum algorithm. We note that, while our architecture does not pose a theoretical upper bound to t_0 , as arbitrarily large detuning energy ϵ could be applied to the empty dots to suppress unwanted processes, very large t_0 will require impractically large voltages on the gates. Similarly, variations in the chemical potential energy $\Delta\mu$ could be overcome by applying an even larger detuning energy ϵ , together by exploiting the regime where the tunnel barriers can be pulsed on and off. However, we require $\Delta\mu < E_C$, where E_C is the charging energy. This significantly reduces overhead in correcting pulses and pulsing amplitude and increases operation speed (see section S1 for details on uniformity and bounds).

Another challenge is to overcome crosstalk, such that physical parameters as ϵ and t_0 can be controlled individually (15). Here, the highly repeatable nature and the presence of only straight lines in our architecture are strongly favorable. Compensating the crosstalk of an individual line by tuning the associated neighbor lines provides a highly symmetric approach. In the following discussion, we assume the presence of such compensation but refer to the main lines only.

Magnetic field layout and ESR

Single-qubit rotations are performed using global ESR striplines (see Fig. 2A) providing in-plane RF magnetic fields (7). A modest external DC magnetic field is applied in the out-of-plane direction.

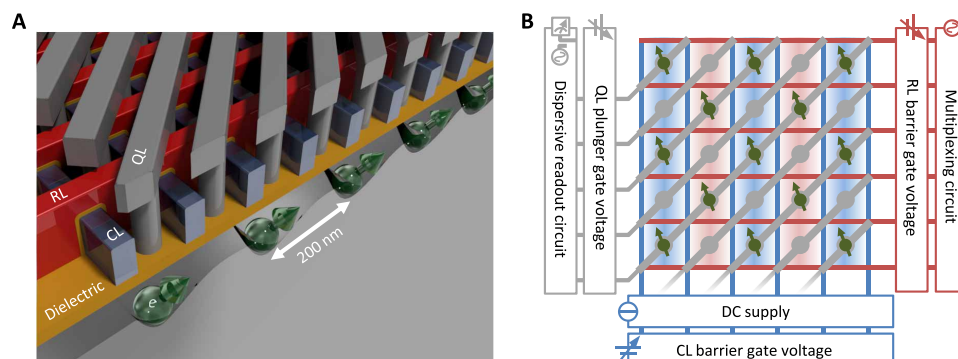


Fig. 1. Design of the quantum dot crossbar array. (A) Three-dimensional model of the array gate structure. The dielectrics in between the various gate layers are left out for clarity. (B) Schematic representation of the 2D quantum dot array. CLs (blue), RLs (red), and QLs (gray) connect the qubit grid to outside electronics for control and readout. A combination of these lines enables qubit selectivity. In the state shown here, half of the quantum dots are occupied with a single electron, where the electron spin encodes the qubit state. The electrons can be shuttled around via the gate voltages, providing a means to couple to nearest neighbors for two-qubit logic gates and readout and to couple to remote qubits for long-range entanglement.

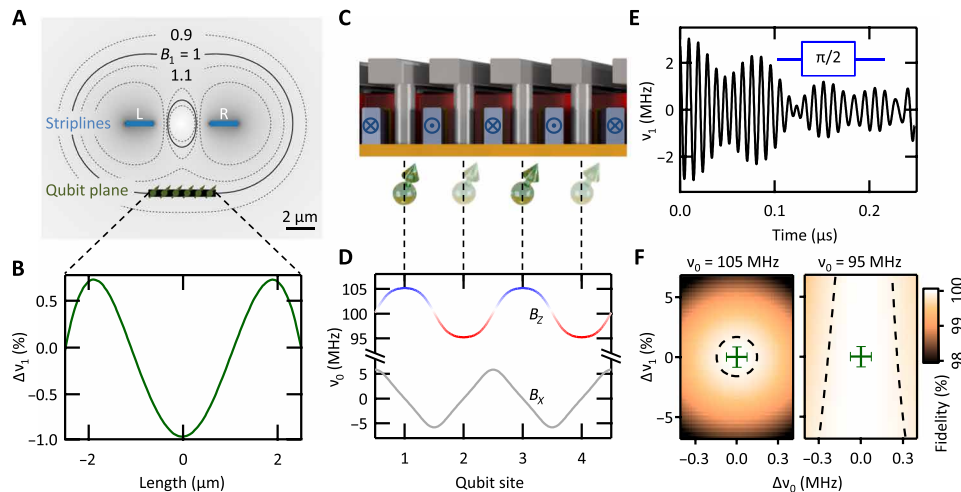


Fig. 2. ESR for single-qubit rotations and magnetic field profile of the crossbar structure. (A) Cross section of the stripline pair (2 μm in width and 6 μm in pitch) positioned 4 μm above the qubit plane. The gray background with black contour lines visualizes the RF magnetic field generated by driving currents through the striplines. (B) The double stripline is optimized to minimize the variations in RF magnetic field at the qubit plane, and we find peak-to-peak values below 2%. (C) Cross section of the qubit plane. A DC alternating in direction between even and odd CLs together with an external out-of-plane field generates a static field that alternates column by column. (D) Corresponding resonance frequency profile. For CLs with dimensions 30 nm wide and 60 nm high and positioned 20 nm above the qubit plane, the required current $I_{\text{CL}} \sim 70 \mu\text{A}$ and current density $j_{\text{CL}} \sim 4 \times 10^{10} \text{ A/m}^2$. The field component B_z has local maxima and minima at the qubit sites, where B_x vanishes, providing qubit addressability and first-order insensitivity to qubit placement. (E) Column-selective qubit pulses are engineered using GRAPE, and here, a selective $\pi/2$ pulse is shown, with fidelities shown in (F). The GRAPE pulse is designed to tolerate static variations in the v_0 and v_1 . The green error bars denote the expected qubit-to-qubit variations, taking into account the design considerations. We conclude that single-qubit rotations can be performed with fidelity higher than 99.9% in the 2D qubit array. Here, we use an electron g -factor of 2 and show the static field by the resonance frequency v_0 and the ESR field by the normal on-resonance Rabi frequency v_1 .

Here, we consider an amplitude of $\sim 3.6 \text{ mT}$, which corresponds to a resonance frequency v_0 of $\sim 100 \text{ MHz}$ for the electron spin. This rather low magnetic field and resonance frequency ease the RF circuit design requirements. In addition, the qubit-to-qubit resonance frequency variation due to spin-orbit coupling (17, 18) is strongly reduced in low magnetic fields and further minimized by applying the magnetic field perpendicular to the interface (19). The ensemble ESR linewidth can then become narrow enough to achieve high-fidelity operation with a global ESR signal. Moreover, we expect improved qubit coherence due to a strongly reduced sensitivity to electrical noise in low fields, as coupling to charge noise via spin-orbit coupling is strongly reduced (19).

Local spin rotations could, in principle, also be implemented by integrating nanomagnets and operation based on electric dipole spin resonance (EDSR) (6). To obtain Rabi frequencies f_{Rabi} beyond 1 MHz, the required transverse field gradient is ~ 0.1 to 0.5 mT/nm for typical driving amplitudes and dot sizes (6). However, while EDSR has proven powerful in single-qubit devices, the integration of nanomagnets in a dense 2D array is much more demanding. In particular, achieving the large required transverse field gradients will also lead to longitudinal field gradients. These will likely affect qubit coherence, shuttling, and two-qubit logic gates. Furthermore, a large gradient appears incompatible with the low-field operation proposed here. Therefore, qubit operation via ESR, requiring minimal field differences, is preferable for spin manipulation in this 2D array design.

To model the striplines and analyze the uniformity and amplitude of the RF fields they generate, we use the Microwave Studio software package from Computer Simulation Technology (CST-MWS) (20). To reach high uniformity across the 2D qubit array, we designed a superconducting stripline pair. We use our CST-MWS model to optimize the relevant dimensions of the stripline design. Furthermore,

to achieve homogeneous fields, the current distribution through the striplines has to be taken into account. For thin-film superconducting striplines, this is, to a large extent, determined by the effective penetration depth λ_{eff} . We find that already for $\lambda_{\text{eff}} > 0.5 \mu\text{m}$, the corresponding RF field inhomogeneity across the 2D array can be less than $\delta v_1 = 2\%$, as shown in Fig. 2B (see section S3 for details). In addition, Rabi driving at 10 MHz requires $\sim 0.6 \text{ mA}$ in each stripline and reasonable current densities $j_{\text{Stripline}} = 3 \times 10^9 \text{ A/m}^2$ in the stripline pair for a thickness of 100 nm.

To achieve qubit addressability, a column-by-column alternating magnetic field is generated by passing DCs with alternating directions through the CL, as shown in Fig. 2 (C and E) (see also section S2). The targeted $\delta v_{\text{CL}} = 10\text{-MHz}$ frequency difference between columns requires current densities $j_{\text{CL}} \sim 4 \times 10^{10} \text{ A/m}^2$ in the gate lines. The integration of superconducting lines as considered here suppresses heat dissipation and minimizes, in addition, potential differences along the lines. The expected field profile along a row of qubits is plotted in Fig. 2D.

Spin-orbit coupling in silicon is strongly enhanced close to an interface and in the presence of large vertical electrical fields, which leads to significant qubit-to-qubit variations in resonance frequency (17, 18, 21). These variations depend on the microscopic interface, and even a single atomic step edge can have a strong impact; it will thus be a significant challenge to overcome these variations by fabrication methods only. In typical silicon metal-oxide-semiconductor quantum dots, the variations in the g -factor are up to $\Delta g/g = 1 \times 10^{-2}$ (17, 21). In SiGe devices, the variations are predicted to be an order of magnitude smaller, $\Delta g/g = 1 \times 10^{-3}$ (18). Possible optimization strategies to reduce variations could focus on the perpendicular electric field or on the material stack. However, by operating in the low magnetic field regime and by applying the field perpendicular to the interface

(19) as proposed here, the qubit-to-qubit variation is expected to vanish, and we take a conservative estimate $\delta v_{\text{SOC}} = 50$ kHz.

Imperfect device fabrication can result in local variations of the magnetic field. This impact is minimized because the magnetic field is self-aligned with the quantum dot barriers defined by the CLs. Furthermore, the magnetic field pattern is designed to have local minima or maxima at the qubit positions, such that the qubit energy splittings are, to first order, insensitive to variations in location. The dominant contributions to variations in v_0 will thus come from variations in the geometry of the gates. For a 1-nm root mean square variation in gate geometry, which can be achieved with current semiconductor manufacturing technology (16), we estimate the corresponding resonance frequency linewidth to be $\delta v_{\text{fab}} = 100$ kHz (see section S2 for more details). On the basis of these considerations, we find a total variation $\delta v_0 = \delta v_{\text{fab}} + \delta v_{\text{SOC}} = 150$ kHz.

For the implementation of global high-fidelity single-qubit operations, it is central that the RF pulses are forgiving with respect to the inhomogeneity in field, as discussed above. At the same time, the pulses need to be highly frequency-selective to ensure that no unintended qubit rotations or phase shifts are induced in the off-resonant columns. Considering these challenges, we applied gradient ascent pulse engineering (GRAPE) for ESR spin control (22), as shown Fig. 2E. With this technique, we can achieve single-qubit fidelities above 99.9% and crosstalk below 0.1% and perform a $\pi/2$ rotation within 250 ns. The tolerance levels for this fidelity are up to 300 kHz in v_0 and more than 3% in v_1 , indicated by the black dashed lines in Fig. 2F. For comparison, we also include (green error bars) the expected qubit-to-qubit variation based on the discussion above, which falls well within the 99.9% fidelity domain. We note that the error bars denote peak-to-peak variations, such that many qubits will have significantly higher fidelity. This implies that further optimization could be done if a certain number of faulty qubits can be tolerated.

Shuttling qubits for addressability and (long-range) entanglement

We now turn to the shuttling of electrons (23, 24) as a means to create addressability for single- and two-qubit logic gates, as well as an efficient method for (remote) qubit swap. The general principle behind the crossbar operation is the combined control of ϵ and t_0 . Since detuning

and tunneling are controlled by different layers of gates, each qubit can be selectively addressed at the corresponding crossing point.

Figure 3 visualizes qubit shuttling along a row or column. Shuttling involves a change in the qubit resonance frequency. Therefore, the electron wave function has to be shifted diabatically with respect to the spin Hamiltonian, so that we can shuttle the qubit between different sites while preserving its spin state. By using a nonlinear pulsing scheme, we can operate the qubit shuttling up to at least 1 GHz with a fidelity higher than 99.9% when accounting for small t_0 and a large pulsing amplitude for uniformity requirements (see section S5).

The difference in the Larmor frequency between adjacent columns can be exploited to construct fast Z-gates operating at 10 MHz (see Fig. 3C). This can be used to correct phase errors or to implement a Z-gate in a quantum algorithm simply by temporarily moving a qubit to an adjacent column for a properly calibrated duration.

Two-qubit logic gates and Pauli spin blockade-based readout

Two types of two-qubit gates can be implemented with quantum dots, namely, the $\sqrt{\text{SWAP}}$ and the controlled-phase (CPHASE) gate (8–10, 25, 26). A direct implementation of the CPHASE gate, however, requires the Zeeman energy difference to be much larger than the exchange coupling, $\delta E_Z \gg J$, to reach high fidelity. The small field gradient $\delta E_Z = 10$ MHz considered here will not fully suppress SWAP-type rotations reducing the fidelity. A possible solution could be to engineer composite pulses, but here, we focus on $\sqrt{\text{SWAP}}$ as the central two-qubit gate (see Fig. 4A). Together with single-qubit rotations, this provides a universal quantum gate set. For example, a CNOT is obtained by interleaving a Z-gate in between two $\sqrt{\text{SWAP}}$ operations, where the Z-gate can be conveniently realized by using the shuttling scheme. To execute the $\sqrt{\text{SWAP}}$, we shuttle two qubits into the same column such that the g-factor difference is minimized, and we tune the qubit exchange by controlling the tunneling barrier gate while keeping the two qubits at the charge symmetry point with the qubit gates (25, 26).

In the low magnetic field regime discussed here, reservoir-based spin initialization and readout are not possible because of thermal broadening. Therefore, we use the Pauli spin blockade between two electrons on neighboring sites for spin initialization and readout. This method has the additional advantage of not requiring a reservoir next

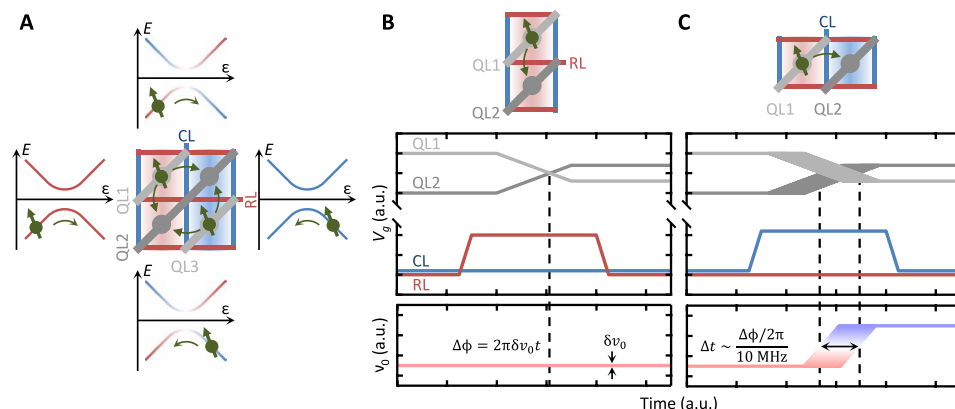


Fig. 3. Qubit shuttling in the crossbar array. (A) By controlling the tunnel coupling and potentials of the dots, qubits can be shuttled around. (B) Shuttling along a column. The sequence consists of setting the tunnel coupling by RL, followed by pulsing the detuning energy. This process leaves the qubit resonance frequency unaffected except for unintended qubit-to-qubit variations. a.u., arbitrary units. (C) Shuttling along a row. This process results in an additional 10-MHz shift in the qubit resonance frequency due to the magnetic field difference between adjacent columns. The shuttling is tuned by controlling the pulsing time on QL1 and QL2. The resulting time delay Δt leads to a controllable phase $\Delta\phi$ applied to the qubit, and this is the basis for our phase updates and Z-gates.

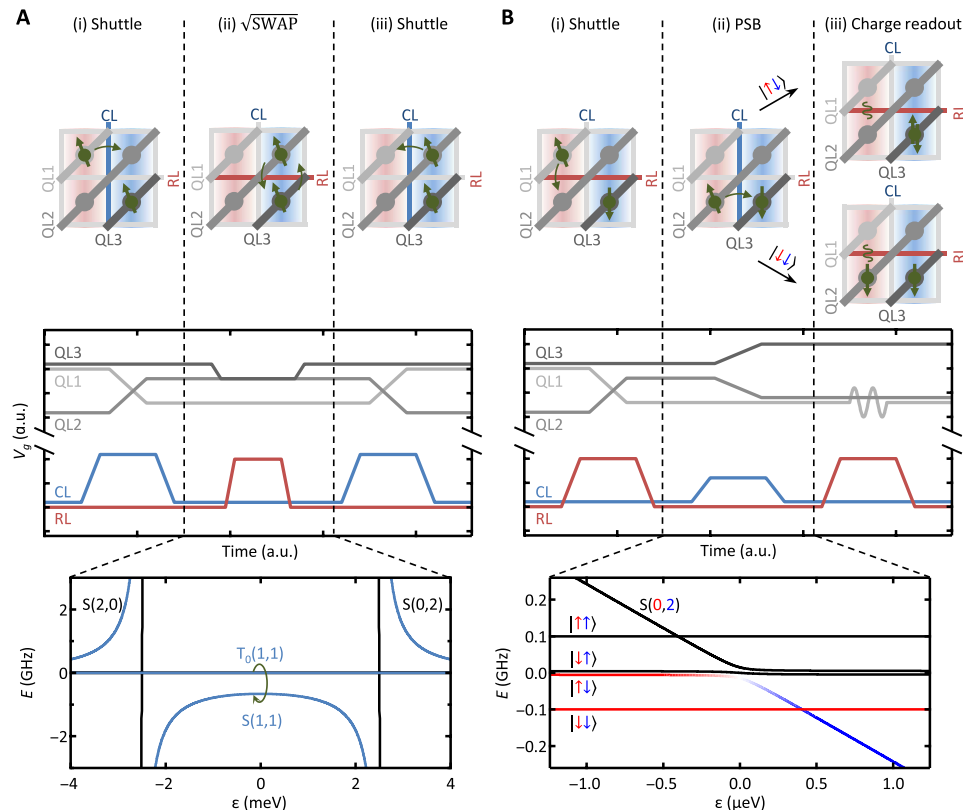


Fig. 4. Two-qubit logic gates and readout. (A) Sequence for $\sqrt{\text{SWAP}}$ gates. By shuttling the respective qubits to the same column, the resonance frequency difference is minimized, enabling a high-fidelity $\sqrt{\text{SWAP}}$. The logic gate is performed at the symmetry point, making the qubits to first-order insensitive to detuning noise, and the interaction is controlled by the associated RL. (B) Spin qubit readout. Here, the respective qubits are shuttled to reside in the same row. The ancillary qubit, located at the blue column with the larger Zeeman energy, is manipulated to the spin-down state. The measurement qubit is adiabatically pulsed. The qubit shuttles when the state is spin up and is blocked when the state is spin down because of the Pauli spin blockade (PSB). Subsequently, the tunnel coupling is turned off, and the charge is locked. Dispersive charge state readout occurs by exploiting an empty neighbor dot.

to the qubit. The protocol relies on the difference in Zeeman energy between the two quantum dots to enable spin parity projection. This difference in energy is created by the same column-by-column alternating magnetic field used to create qubit addressability, and readout is performed between neighboring quantum dots in different columns.

The Pauli spin blockade spin-to-charge conversion scheme is plotted in Fig. 4B. Instead of shuttling along a row, which brings two qubits to adjacent sites in the same column (same resonance frequency), the qubit is now moved along a column. This brings it next to a qubit in a different column, providing the difference in Zeeman energy that is necessary for readout. In the sequence shown in Fig. 4B, the qubit with the smaller Zeeman energy (red background) will be read out. The qubit with the larger Zeeman energy (blue background) serves as an ancillary qubit and must be in the spin-down state, so that other triplet states [see black lines in Fig. 4B (bottom)] can be neglected. If required, single-qubit pulses could be applied to manipulate the ancillary qubit to the spin-down state. By tuning to the configuration where the singlet state becomes the ground state on the ancillary dot, the Pauli spin blockade will prevent (allow) the spin-down (up) electron to shuttle to the ancillary qubit. The above process completes the spin-to-charge conversion, and the spin state can be inferred from the charge occupation. A conversion fidelity higher than 99.9% can be achieved with a 3-MHz gate pulsing speed (see section S5). We note that, in another protocol, the ancillary qubit can be in the

spin-up state, provided that it resides in the column with the smaller Zeeman energy (see section S6). This possibility could prove to be powerful in quantum error correction cycles, as it avoids the need to actively correct errors. In addition, the reverse of the Pauli spin blockade spin-to-charge conversion pulsing process is used for qubit initialization. In the scheme shown in Fig. 4B, if the Pauli spin blockade prevented the qubit to move to the ancillary qubit in the readout step, then it is and remains in the spin-down state. If the qubit moved to the ancillary qubit, it was in the spin-up state before readout. After detuning back, it will return to the spin-up state. In both cases, the ancillary qubit will end up in the spin-down state.

Directly after the Pauli spin blockade spin-to-charge conversion, we switch off the interdot tunnel coupling with CL so that the charge state is disconnected from the spin configuration. In this mode, the state is not sensitive to spin relaxation, thereby increasing the readout fidelity. This can be exploited for delayed readout schemes, such as charge sensor-based readout, by shuttling to the periphery of the 2D array. However, here, we consider gate-based dispersive readout (13, 27) for an on-site readout of the charge state, as shown in Fig. 4B. By applying an RF carrier signal to the qubit gates and coupling the dot to an adjacent empty dot, the charge state can be extracted from the dispersive signal. When there is charge occupation, the interdot oscillation driven by the RF carrier gives an additional quantum capacitance, leading to a different reflected signal compared to the state without

charge occupation. By measuring the reflected signal, we thus determine the qubit state.

Parallel operation

For an efficient quantum computing scheme, simultaneous operation is essential. Here, we discuss how the local operations introduced above can be advanced toward line-by-line or even near-global operation. Contrary to local operations, parallel operations result in active gates crossing at quantum dots that are not targeted (see Fig. 5). This may lead to undesired operations. However, these can be prevented by selectively occupying the quantum dots and specific control of ϵ or t_0 , such that away from the targeted locations signals are only applied to empty quantum dots or to quantum dots with empty neighbors. We also note that the crossbar control scheme could affect non-targeted qubits under the active gates, for example, via Stark shifts. Understanding and managing the consequences of these errors are thus highly important. In our architecture, we minimize these errors by operating at low magnetic field, designing qubit columns with well-separated resonance frequencies, and using adiabatic pulsing schemes, such that the crosstalk errors are significantly smaller than the errors on the targeted qubits. Furthermore, we note that these errors can be largely corrected in subsequent operations with a manage-

able overhead, for example, by implementing an additional phase-controlled shuttling step.

Figure 5A shows an example of a line-by-line operation of controlled-phase shuttling. To properly control the timing, it is crucial to individually pulse the QL. Still, parallel shuttling operations can be implemented along one column or row, enabled by lifting the barriers controlled by one CL or RL, respectively. These CLs can be time-controlled individually to correct the qubit-to-qubit variations, such that the shuttled qubits have the correct phase after the shuttling. The line-by-line shuttle can be performed within 1 ns with a fidelity beyond 99.9% (see section S5).

An approach to performing simultaneous two-qubit logic operations on the qubit module could be to shuttle line by line all target qubits to the associated control qubits and then perform $\sqrt{\text{SWAP}}$ operations line by line. However, this will lead to qubit configurations where targeted qubits share gate lines disabling individual gate control, which is essential for high-fidelity operation. To overcome this, we propose sequences whereby a single column (or row) of qubits is shuttled first, followed by the desired operation and shuttle back, and then the sequence is continued by operating the next line of qubits of the module until all qubits are addressed. This protocol is demonstrated in Fig. 5B, which shows the configuration after shuttling a single

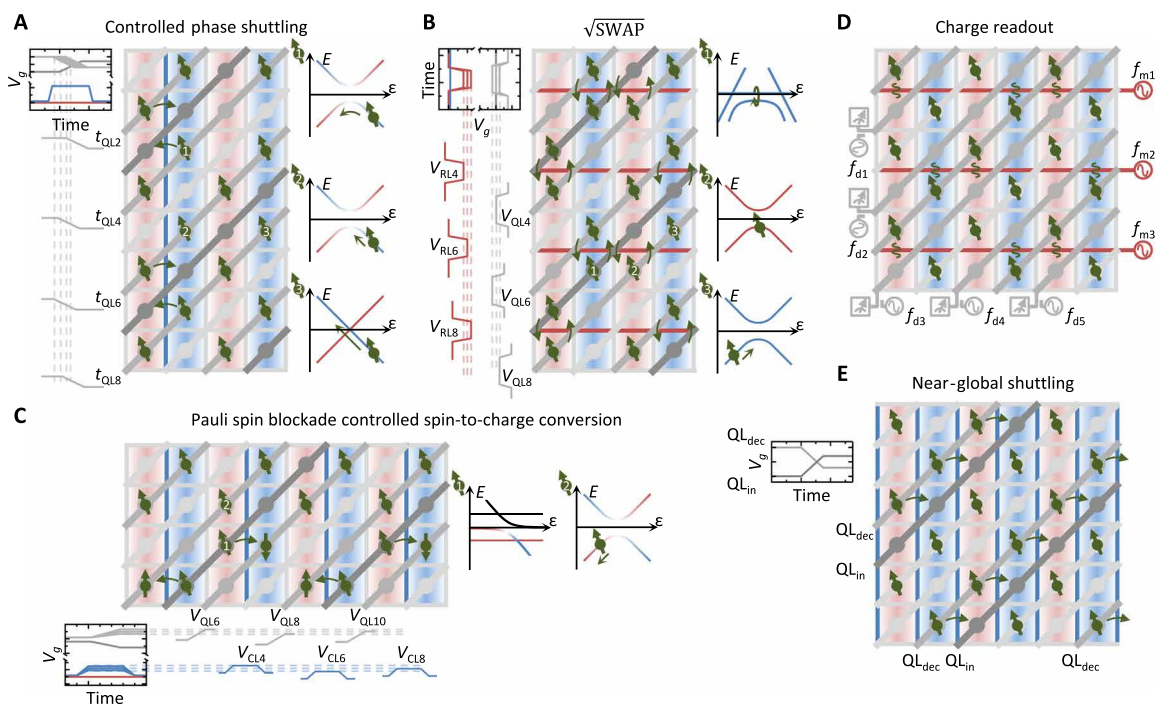


Fig. 5. Parallel operation. (A to C) Simultaneous operation of controlled-phase shuttling (A), two-qubit $\sqrt{\text{SWAP}}$ operations (B), and spin-to-charge conversion (C) can be achieved in a line-by-line manner. In each figure, inset (1) denotes the energy-detuning diagram of the targeted qubit(s). Insets (2) and (3) show the consequence on the remaining qubits, where detuning, tunnel coupling, or the local magnetic field minimizes errors. (D and E) Shuttling without phase control (D) and charge readout (E) can be performed in a near-global manner. (A) Shuttling of qubits. Parallelism is obtained along one direction, and tunability is obtained along another direction, and the respective gates control the timing and detuning to overcome qubit-to-qubit variations. Here, the target qubits shuttle from column to column, whereas the other qubits are blocked by ϵ or t_0 . (B) Two-qubit logic gates. $\sqrt{\text{SWAP}}$ operations only occur between tunnel-coupled neighboring qubits. The remaining qubits do not interact but could shuttle in a column. The resulting (small) phase shift can be corrected by the consecutive shuttle event in the line-by-line operation. In (C), Pauli spin blockade spin-to-charge conversion occurs between tunnel-coupled qubits. Qubits coupled to an empty dot do not shuttle, prevented by the energy alignment, since we require $\Delta\mu < E_C$. (D) Shuttling without phase control enables to construct a variety of shuttle patterns that can be operated almost globally; the schematic here shows the simultaneous shuttling of half of the qubits one site to the right. (E) The dispersive charge readout, performed after the spin-to-charge conversion shown in (C), can be performed simultaneously by including frequency multiplexing. The RF carrier on QL (f_d) is then modulated by the application of additional multiplexing RF pulses (f_m) to RL.

column of qubits. To overcome variations in tunnel couplings and chemical potentials, we tune the amplitude and duration of the pulses applied to the respective RLs and CLs individually for each two-qubit gate to achieve the desired operation. For example, operations can be performed at the detuning-noise insensitive charge symmetry point (25, 26). Consequently, the line-by-line control does not limit the operation speed, and we envision operation frequencies in the range of 10 to 100 MHz for two-qubit logic gates.

Simultaneous readout consists of a spin-to-charge conversion step and charge readout step. First, a row of qubits is shuttled, resulting in the configuration shown in Fig. 5C. After that, the parameters ϵ and t_0 can be individually controlled to convert spin to charge. In this specific sequence here, qubits are alternately shuttled up or down along the row, which leads to a configuration that is typically compatible with error correction sequences (2, 28). However, there may be instances where a different configuration is required, and this could reduce the spin-to-charge conversion to half the speed compared to line-by-line.

Global charge readout requires us to distinguish between qubits connected to the same QL. This is achieved via frequency multiplexing, as shown in Fig. 5D. Here, an additional voltage modulation is applied to the RL. The separation of spin-to-charge conversion and charge readout in different steps has a particular advantage. While the initial spin-to-charge conversion must be performed line-by-line, it can be done relatively fast. The readout of charge is likely slower, and to overcome the nonuniformity in $\Delta\mu$, a large detuning has to be applied. Instead of a single-step readout, we sequentially read out for different detuning and group the qubits according to their detuning (see section S5). This sequential readout, as compared to the line-by-line approach, has the advantage that is independent on the number of qubits and will be efficient for large qubit modules. The total readout time will strongly depend on the performance of dispersive readout at the single-qubit level, now under intensive research. However, the protocol here shows that the slowdown with increasing numbers of qubits can be controlled.

Near-global operation is possible when phase control is not required. This may have multiple applications, for example, in achieving long-range coupling. In these protocols, multiple shuttles can be performed with a single phase match at the start or at an arbitrary point. An example of global shuttling is shown in Fig. 5E, where half of the qubits are simultaneously moved. Shuttling requires adiabatic movement with respect to the tunnel coupling, and the demand is most stringent close to the anticrossing point. Because of the qubit-to-qubit variations, it may not be possible to go beyond a linear detuning pulse, as each pair can have the anticrossing at a different location. This consequently limits the shuttle speed. Nonetheless, for a $\Delta\mu = 2$ meV, shuttling can be at a 1-GHz rate when $t_0 > 25$ GHz (see section S5). This simultaneous shuttling can be highly important for advanced error correction codes that require long-distance coupling, such as the 3D gauge color code (29).

A network of qubit modules

For truly large-scale quantum computation, we envision a network consisting of a large number of interconnected qubit modules. While the layout of such an architecture crucially depends on the specific qubit module implementation and therefore goes beyond the current proposal, a possible repeatable tile is depicted in Fig. 6 (see also section S7). In addition to the central array hosting the qubit module, the quantum dot grid is extended in a simpler structure consisting of barrier gates only, thereby strongly reducing the number of re-

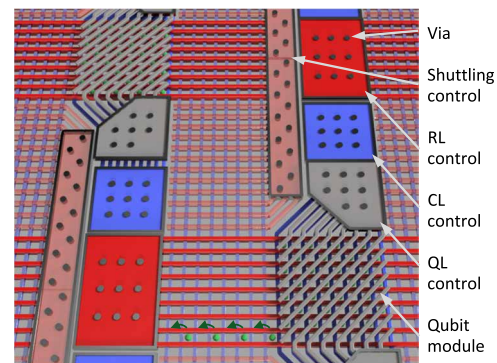


Fig. 6. Prospects for connecting qubit modules. Individual qubit modules (targeted to be of order 1000 qubits; for clarity reasons, smaller modules are shown here) are connected together using long-range shuttle highways. The parallelism of the long-range shuttlers strongly reduces the number of control lines, providing space to integrate local electronics or vertical vias to interconnect the qubit array to outside electronics. Individual qubit modules could be operated using specific codes or be programmed to host, for example, a single logical qubit.

quired control lines. These shuttling dots cannot be fully controlled but allow for the transportation of qubits (23, 24). With this approach, qubit modules can then be connected together, where the available space can be used for local electronics (4, 11, 13) or wiring fan-out. Transportation of, for example, a column of qubits from the edge of one module to another module would then provide a large range of possibilities for quantum algorithms, since it would create a large virtual array of coupled qubits with a certain degree of long-range coupling.

DISCUSSION

One of the greatest challenges in the area of scalability is avoiding an interconnect bottleneck. Here, we have proposed a scalable solution for spin qubits based on crossbar technology. While this technology limits control, we have developed general operation schemes based on partial sequential control. The increased operation time due to sequential control is warranted by the very long coherence times of quantum dot spin qubits, with experimental demonstrations already up to 28 ms (7). We have shown operation schemes for phase-controlled shuttling, two-qubit logic gates, and spin-to-charge conversion. These operations can have a targeted execution time well below 1 μ s. The resulting loss of coherence due to the waiting time when operating in a line-by-line manner could be well below 10^{-3} in a 1000-qubit module using suitable echo sequences. The shuttling proposed here can be performed simultaneously within 1 ns, enabling even more than 10^7 operations, and could provide an excellent method to create long-range entanglement or remote qubit SWAP. Readout could become fast by global operation, and measurement-free quantum error correction schemes could reduce the need for frequent readout (30, 31).

While the proposed structure is compatible with existing technology, several aspects of the design require an experimental demonstration. Studies of the uniformity level on extended quantum dot arrays will validate the shared gate control scheme. Spin qubit operation in moderate magnetic fields have been demonstrated (32), but more work is needed to investigate the limits of single-qubit operation fidelity. After encouraging results of shuttling electron spin states in GaAs quantum dots (24), these experiments should be repeated in Si and, in particular, in ^{28}Si to investigate the fidelity that can be

reached for the coherent shuttling, as proposed here. Studies of the fidelity of dispersive charge detection will enable to compare simultaneous readout with alternative shuttling and serial detection by one or several fast charge sensors. A feasibility study of quantum error correction (28) on this architecture on the single logical qubit scale suggests extremely low error rates, while the global shuttling scheme promises avenues to incorporate multiple logical qubits in a single module. Further scaling going beyond these modules will introduce new challenges, and interfacing protocols should account for the extra elements and the accompanying errors. In particular, long-range coupling of qubits will be crucial for ultimate scaling (11), and the errors due to the limited control inside these couplers should be included in future error analysis. If advances in qubit control continue to improve and lead to all fidelities higher than 99.9%, then the architecture discussed here provides an excellent way forward to large-scale quantum computation.

The proposed architecture supports universal quantum computation in a fault-tolerant manner (28), where the ability to shuttle qubits over large distances in principle provides means to realize quantum error correction schemes and quantum circuit implementations otherwise reserved for nonplanar architectures. Within one qubit module, the highly flexible nature of the presented architecture makes it amenable to the use of a variety of topological error correction codes (28). For planar codes, this includes the surface code (2), which has a fault tolerance threshold as high as 1% (33) and moreover can be implemented using entangling gates between qubits that are adjacent on a 2D surface. A distance-three surface code would fit in a 7×7 quantum dot module, and a successful implementation would present a milestone on the path toward fault-tolerant quantum computation.

The proposed architecture is also amenable to other 2D local topological error codes such as the 2D color code (29), which has a lower threshold (34) but supports a more expansive set of logical operations. Finally, we also envision that the use of qubit shuttling will enable the implementation of error correction schemes requiring long-range entanglement such as the 3D gauge color code (29). This approach has several highly desirable properties including low stabilizer generator weight, the possibility of a high error threshold (29), and the ability to perform (through a procedure called gauge fixing) a universal gate set in a fault-tolerant manner. This last property would preclude the need for procedures such as magic-state distillation, which are currently foreseen to take up the vast majority of computing resources in other fault-tolerant quantum computation schemes.

We remark that entangling operations of surface code logical qubits encoded in two different qubit modules can be performed by shuttling only the qubits at the edges to the other qubit modules (see Fig. 6) (35) and subsequently returning the qubits to the original module. This avoids the necessity to shuttle all qubits in one module to the next to perform two-qubit gates between logical qubits.

We could foresee lower performance regimes or faulty qubits on the chip, for example, due to the qubit-to-qubit variation induced by the ESR stripline pair. One way to address faulty sites within one qubit module would be to change the actual quantum error-correcting code to encode one (or more) logical qubits with fewer physical qubits using the remaining qubits in the vicinity (35). Yet, it is clear that this introduces inhomogeneity in the classical control requirements of the individual modules and greatly complicates two-qubit gates between two logical qubits, as they are now encoded using different codes.

Depending on the fidelity of the long-distance shuttling operations in the fabricated devices, however, another path could be to turn off qubit modules completely if the noise exceeds a certain threshold. As a consequence, we may need to shuttle qubits over longer distances to perform two-qubit operations on logical qubits but would have the ability to select the desired good qubit modules. This is particularly promising in this architecture, given the ability to shuttle fast and with high fidelity.

A particular challenge is to map quantum circuits to our architecture. For this, a variety of classical methods exist. To gain maximum advantage of the ability to shuttle qubits, the long-distance shuttling operations are ideally fast compared to general gate speeds. In this case, the architecture becomes virtually nonplanar, which can yield significant savings in overhead (36).

While many traditional quantum algorithms such as Shor's factoring algorithm require a large number of qubits, few-qubit applications are slowly beginning to emerge. In recent years, interest in electronic structure quantum simulation has culminated in small-scale experimental implementations (37, 38). Larger simulation algorithms will have to deal with entangling large amounts of qubits along certain paths across the device, as introduced by the standard mapping of second quantization. The switching to different mappings, on the other hand (39), reduces the amount of gates but does not solve the connectivity problems, which the proposed architecture is a promising candidate to tackle. Shuttling and the native $\sqrt{\text{SWAP}}$ gates might also be used to move certain auxiliary qubits around, which allows for significant decreases in the depth of the resulting quantum circuit (40).

Upscaling toward the numbers of qubits required for these algorithms, including few qubit applications, represents a formidable challenge. However, we envision that the proposed architecture based on shared control and flexible qubit shuttling can provide a unique shortcut toward large-scale quantum computation.

MATERIALS AND METHODS

The inhomogeneity of the ESR stripline (Fig. 2B) was simulated by a 3D model created with the CST-MWS (20). More details and discussions on the stripline parameters can be found in section S3.

The column-by-column alternating spin resonance frequency (Fig. 2D) was estimated by calculating the magnetic field generated from DCs through CLs with alternating directions. The field was calculated with the Biot-Savart law and then converted to resonance frequency using a g -factor of 2. More details and discussions on the effect of fabrication errors can be found in section S2.

The column-selective qubit pulse (Fig. 2E) was generated using gradient ascent pulse engineering (GRAPE) (22). For tolerance to qubit nonuniformity, we optimize the average fidelity on four qubits: two qubits with $\nu_0 = 105 \pm 0.1$ MHz for a $\pi/2$ rotation and two qubits with $\nu_0 = 95 \pm 0.1$ MHz for a null operation. More details and discussions on the comparison to square ESR pulse can be found in section S4.

The qubit shuttling operation and Pauli spin blockade spin-to-charge conversion were simulated numerically with the time evolution of the quantum states. Figure S5 shows that over a large detuning and tunnel coupling range, qubits can be pulsed to the desired states with higher than 99.9% fidelity. More details can be found in section S5.

SUPPLEMENTARY MATERIALS

Supplementary material for this article is available at <http://advances.sciencemag.org/cgi/content/full/4/7/eaar3960/DC1>

Section S1. Tolerance to quantum dot inhomogeneity
 Section S2. Column-by-column alternating static magnetic field
 Section S3. Inhomogeneity of the ESR stripline
 Section S4. GRAPE pulse for spin rotation
 Section S5. Shuttling fidelity
 Section S6. Pauli spin blockade spin-to-charge conversion with ancillary qubit in the spin-up state
 Section S7. Shuttling bus for a 2D array module
 Fig. S1. Impact of misalignment and errors in gate and dot dimensions.
 Fig. S2. Overall resonance frequency error as a function of fabrication error.
 Fig. S3. Stripline schematic and simulation results.
 Fig. S4. GRAPE pulse optimization for high-fidelity single-qubit gates.
 Fig. S5. Charge shuttling process.
 Fig. S6. Scheme for Pauli spin blockade spin-to-charge conversion with ancillary qubit in the spin-up state.
 Fig. S7. Connecting qubit modules.
 References (41–44)

REFERENCES AND NOTES

- D. Loss, D. P. DiVincenzo, Quantum computation with quantum dots. *Phys. Rev. A* **57**, 120–126 (1998).
- A. G. Fowler, M. Mariantoni, J. M. Martinis, A. N. Cleland, Surface codes: Towards practical large-scale quantum computation. *Phys. Rev. A* **86**, 032324 (2012).
- C. D. Hill, E. Peretz, S. J. Hile, M. G. House, M. Fuechsle, S. Rogge, M. Y. Simmons, L. C. L. Hollenberg, A surface code quantum computer in silicon. *Sci. Adv.* **1**, e1500707 (2015).
- M. Veldhorst, H. G. J. Eenink, C. H. Yang, A. S. Dzurak, Silicon CMOS architecture for a spin-based quantum computer. *Nat. Commun.* **8**, 1766 (2017).
- B. Lekitsch, S. Weidt, A. G. Fowler, K. Mølmer, S. J. Devitt, C. Wunderlich, W. K. Hensinger, Blueprint for a microwave trapped ion quantum computer. *Sci. Adv.* **3**, e1601540 (2017).
- E. Kawakami, P. Scarlino, D. R. Ward, F. R. Braakman, D. E. Savage, M. G. Lagally, M. Friesen, S. N. Coppersmith, M. A. Eriksson, L. M. K. Vandersypen, Electrical control of a long-lived spin qubit in a Si/SiGe quantum dot. *Nat. Nanotechnol.* **9**, 666–670 (2014).
- M. Veldhorst, J. C. C. Hwang, C. H. Yang, A. W. Leenstra, B. de Ronde, J. P. Dehollain, J. T. Muhonen, F. E. Hudson, K. M. Itoh, A. Morello, A. S. Dzurak, An addressable quantum dot qubit with fault-tolerant control-fidelity. *Nat. Nanotechnol.* **9**, 981–985 (2014).
- M. Veldhorst, C. H. Yang, J. C. C. Hwang, W. Huang, J. P. Dehollain, J. T. Muhonen, S. Simmons, A. Laucht, F. E. Hudson, K. M. Itoh, A. Morello, A. S. Dzurak, A two-qubit logic gate in silicon. *Nature* **526**, 410–414 (2015).
- T. F. Watson, S. G. J. Philips, E. Kawakami, D. R. Ward, P. Scarlino, M. Veldhorst, D. E. Savage, M. G. Lagally, M. Friesen, S. N. Coppersmith, M. A. Eriksson, L. M. K. Vandersypen, A programmable two-qubit quantum processor in silicon. *Nature* **555**, 633–637 (2018).
- D. M. Zajac, A. J. Sigillito, M. Russ, F. Borjans, J. M. Taylor, G. Burkard, J. R. Petta, Resonantly driven CNOT gate for electron spins. *Science* **359**, 439–442 (2018).
- L. M. K. Vandersypen, H. Bluhm, J. S. Clarke, A. S. Dzurak, R. Ishihara, A. Morello, D. J. Reilly, L. R. Schreiber, M. Veldhorst, Interfacing spin qubits in quantum dots and donors—Hot, dense and coherent. *Npj Quantum Inf.* **3**, 34 (2017).
- B. Stackhouse, S. Bhimji, C. Bostak, D. Bradley, B. Cherkauer, J. Desai, E. Francom, M. Gowan, P. Gronowski, D. Krueger, C. Morganti, A 65 nm 2-billion transistor quad-core Itanium processor. *IEEE J. Solid State Circuits* **44**, 18–31 (2009).
- S. Schaal, S. Barraud, J. J. L. Morton, M. F. Gonzalez-Zalba, Conditional dispersive readout of a CMOS quantum dot via an integrated transistor circuit. *Phys. Rev. Appl.* **9**, 054016 (2017).
- M. G. Borselli, K. Eng, R. S. Ross, T. M. Hazard, K. S. Holabird, B. Huang, A. A. Kiselev, P. W. Deelman, L. D. Warren, I. Milosavljevic, A. E. Schmitz, M. Sokolich, M. F. Gyure, A. T. Hunter, Undoped accumulation-mode Si/SiGe quantum dots. *Nanotechnology* **26**, 375202 (2015).
- T. Hensinger, T. Fujita, L. Janssen, X. Li, C. J. Van Diepen, C. Reichl, W. Wegscheider, S. Das Sarma, L. M. K. Vandersypen, Quantum simulation of a Fermi–Hubbard model using a semiconductor quantum dot array. *Nature* **548**, 70–73 (2017).
- W. van der Zande, EUVL exposure tools for HVM: It's under (and about) control, in *Proceedings of the EUV and Soft X-ray Source Workshop*, Amsterdam, Netherlands, 7 to 9 November 2016.
- M. Veldhorst, R. Ruskov, C. H. Yang, J. C. C. Hwang, F. E. Hudson, M. E. Flatté, C. Tahan, K. M. Itoh, A. Morello, A. S. Dzurak, Spin-orbit coupling and operation of multivalley spin qubits. *Phys. Rev. B* **92**, 201401 (2015).
- R. Ferdous, E. Kawakami, P. Scarlino, M. P. Nowak, D. R. Ward, D. E. Savage, M. G. Lagally, S. N. Coppersmith, M. Friesen, M. A. Eriksson, L. M. K. Vandersypen, R. Rahman, Valley dependent anisotropic spin splitting in silicon quantum dots. arXiv:1702.06210 (2017).
- R. M. Jock, N. T. Jacobson, P. Harvey-Collard, A. M. Mounce, V. Srinivasa, D. R. Ward, J. Anderson, R. Manginell, J. R. Wendt, M. Rudolph, T. Pluym, J. K. Gamble, A. D. Baczewski, W. M. Witzel, M. S. Carroll, A silicon metal-oxide-semiconductor electron spin-orbit qubit. *Nat. Commun.* **9**, 1768 (2018).
- CST MICROWAVE STUDIO®, CST Computer Simulation Technology AG, www.cst.com.
- R. Ferdous, K. W. Chan, M. Veldhorst, J. C. C. Hwang, C. H. Yang, G. Klimeck, A. Morello, A. S. Dzurak, R. Rahman, Interface induced spin-orbit interaction in silicon quantum dots and prospects for scalability. arXiv:1703.03840 (2017).
- N. Khaneja, T. Reiss, C. Kehlet, T. Schulte-Herbrüggen, S. J. Glaser, Optimal control of coupled spin dynamics: Design of NMR pulse sequences by gradient ascent algorithms. *J. Magn. Reson.* **172**, 296–305 (2005).
- J. M. Taylor, H.-A. Engel, W. Dür, A. Yacoby, C. M. Marcus, P. Zoller, M. D. Lukin, Fault-tolerant architecture for quantum computation using electrically controlled semiconductor spins. *Nat. Phys.* **1**, 177–183 (2005).
- T. A. Baart, M. Shafiei, T. Fujita, C. Reichl, W. Wegscheider, L. M. K. Vandersypen, Single-spin CCD. *Nat. Nanotechnol.* **11**, 330–334 (2016).
- M. D. Reed, B. M. Maune, R. W. Andrews, M. G. Borselli, K. Eng, M. P. Jura, A. A. Kiselev, T. D. Ladd, S. T. Merkel, I. Milosavljevic, E. J. Pritchett, M. T. Rakher, R. S. Ross, A. E. Schmitz, A. Smith, J. A. Wright, M. F. Gyure, A. T. Hunter, Reduced sensitivity to charge noise in semiconductor spin qubits via symmetric operation. *Phys. Rev. Lett.* **116**, 110402 (2016).
- F. Martins, F. K. Malinowski, P. D. Nissen, E. Barnes, S. Fallahi, G. C. Gardner, M. J. Manfra, C. M. Marcus, F. Kuemmeth, Noise suppression using symmetric exchange gates in spin qubits. *Phys. Rev. Lett.* **116**, 116801 (2016).
- J. I. Colless, A. C. Mahoney, J. M. Hornibrook, A. C. Doherty, D. J. Reilly, H. Lu, A. C. Gossard, Dispersive readout of a few-electron double quantum dot with fast rf gate sensors. *Phys. Rev. Lett.* **110**, 046805 (2013).
- J. Helsen, M. Steudtner, M. Veldhorst, S. Wehner, Quantum error correction in crossbar architectures. *Quantum Sci. Technol.* **3**, 3 (2018).
- H. Bombin, Gauge color codes: Optimal transversal gates and gauge fixing in topological stabilizer codes. *New J. Phys.* **17**, 083002 (2015).
- V. Nebendahl, H. Häffner, C. F. Roos, Optimal control of entangling operations for trapped-ion quantum computing. *Phys. Rev. A* **79**, 012312 (2009).
- G. A. Paz-Silva, G. K. Brennen, J. Twamley, Fault tolerance with noisy and slow measurements and preparation. *Phys. Rev. Lett.* **105**, 100501 (2010).
- M. A. Fogarty, K. W. Chan, B. Hensen, W. Huang, T. Tanttu, C. H. Yang, A. Laucht, M. Veldhorst, F. E. Hudson, K. M. Itoh, D. Culcer, A. Morello, A. S. Dzurak, Integrated silicon qubit platform with single-spin addressability, exchange control and robust single-shot singlet-triplet readout. arXiv:1708.03445 (2017).
- D. S. Wang, A. G. Fowler, L. C. L. Hollenberg, Surface code quantum computing with error rates over 1%. *Phys. Rev. A* **83**, 020302 (2011).
- A. J. Landahl, J. T. Anderson, P. R. Rice, Fault-tolerant quantum computing with color codes. arXiv:1108.5738 (2011).
- C. Horsman, A. G. Fowler, S. Devitt, R. V. Meter, Surface code quantum computing by lattice surgery. *New J. Phys.* **14**, 123011 (2012).
- R. Beals, S. Brierley, O. Gray, A. Harrow, S. Kutin, N. Linden, D. Shepherd, M. Stather, Efficient distributed quantum computing. *Proc. R. Soc. A* **469**, 20120686 (2013).
- P. J. J. O'Malley, R. Babbush, I. D. Kivlichan, J. Romero, J. R. McClean, R. Barends, J. Kelly, P. Roushan, A. Tranter, N. Ding, B. Campbell, Y. Chen, Z. Chen, B. Chiaro, A. Dunsworth, A. G. Fowler, E. Jeffrey, A. Megrant, J. Y. Mutus, C. Neill, C. Quintana, D. Sank, A. Vainsencher, J. Wenner, T. C. White, P. V. Coveney, P. J. Love, H. Neven, A. Aspuru-Guzik, J. M. Martinis, Scalable quantum simulation of molecular energies. *Phys. Rev. X* **6**, 031007 (2016).
- A. Kandala, A. Mezzacapo, K. Temme, M. Takita, M. Brink, J. M. Chow, J. M. Gambetta, Hardware-efficient quantum optimizer for small molecules and quantum magnets. *Nature* **549**, 242–246 (2017).
- S. B. Bravyi, A. Y. Kitaev, Fermionic quantum computation. *Ann. Phys.* **298**, 210–226 (2002).
- M. B. Hastings, D. Wecker, B. Bauer, M. Troyer, Improving quantum algorithms for quantum chemistry. *Quantum Inf. Comput.* **15**, 1–21 (2015).
- H. Bartolf, A. Engel, A. Schilling, K. Il'in, M. Siegel, H.-W. Hübers, A. Semenov, Current-assisted thermally activated flux liberation in ultrathin nanopatterned NbN superconducting meander structures. *Phys. Rev. B* **81**, 024502 (2010).
- A. I. Gubin, K. S. Il'in, S. A. Vitusevich, M. Siegel, N. Klein, Dependence of magnetic penetration depth on the thickness of superconducting Nb thin films. *Phys. Rev. B* **72**, 064503 (2005).
- X. Li, E. Barnes, J. P. Kestner, S. Das Sarma, Intrinsic errors in transporting a single-spin qubit through a double quantum dot. *Phys. Rev. A* **96**, 012309 (2017).
- R. Mizuta, R. M. Otxoa, A. C. Betz, M. F. Gonzalez-Zalba, Quantum and tunneling capacitance in charge and spin qubits. *Phys. Rev. B* **95**, 045414 (2017).

Acknowledgments

Funding: M.V. acknowledges support from the Netherlands Organisation of Scientific Research (NWO) Vidi program. J.H. and S.W. are funded by an NWO Vidi grant, a European Research Council (ERC) Starting Grant, and STW Netherlands. M.S. is funded by NWO/OCW and an ERC synergy grant. L.M.K.V. acknowledges support from the NWO Vici program. **Author contributions:** All authors contributed to the architecture design. R.L. and M.V. developed the physical layout and the quantum operations. L.P. simulated the GRAPE pulses. J.P.D. performed the CST-MWS simulations. R.L., D.P.F., and M.V. wrote the manuscript with input from all authors. **Competing interests:** L.M.K.V., K.J.S., M.V., and J.S.C. are inventors on a patent application related to this work filed by Intel Corporation (application no. PCT/US2017/039155; filing date, 24 June 2017). The other authors declare that they have no competing interests. **Data and materials availability:**

All data needed to evaluate the conclusions in the paper are present in the paper and/or the Supplementary Materials. Additional data related to this paper may be requested from the authors.

Submitted 8 November 2017

Accepted 29 May 2018

Published 6 July 2018

10.1126/sciadv.aar3960

Citation: Li, L., Petit, D. P., Franke, J. P., Dehollain, J., Helsen, M., Steudtner, N. K., Thomas, Z. R., Yoscovits, K. J., Singh, S., Wehner, L. M. K., Vandersypen, J. S., Clarke, M., Veldhorst, A crossbar network for silicon quantum dot qubits. *Sci. Adv.* **4**, eaar3960 (2018).

A crossbar network for silicon quantum dot qubits

Ruoyu Li, Luca Petit, David P. Franke, Juan Pablo Dehollain, Jonas Helsen, Mark Steudtner, Nicole K. Thomas, Zachary R. Yoscovits, Kanwal J. Singh, Stephanie Wehner, Lieven M. K. Vandersypen, James S. Clarke and Menno Veldhorst

Sci Adv 4 (7), eaar3960.
DOI: 10.1126/sciadv.aar3960

ARTICLE TOOLS

<http://advances.sciencemag.org/content/4/7/eaar3960>

SUPPLEMENTARY MATERIALS

<http://advances.sciencemag.org/content/suppl/2018/07/02/4.7.eaar3960.DC1>

REFERENCES

This article cites 38 articles, 4 of which you can access for free
<http://advances.sciencemag.org/content/4/7/eaar3960#BIBL>

PERMISSIONS

<http://www.sciencemag.org/help/reprints-and-permissions>

Use of this article is subject to the [Terms of Service](#)

Science Advances (ISSN 2375-2548) is published by the American Association for the Advancement of Science, 1200 New York Avenue NW, Washington, DC 20005. 2017 © The Authors, some rights reserved; exclusive licensee American Association for the Advancement of Science. No claim to original U.S. Government Works. The title *Science Advances* is a registered trademark of AAAS.