

Scalable parity architecture with a shuttling-based spin qubit processor

Florian Ginzel¹,¹ Michael Fellner^{2,3},^{2,3} Christian Ertler¹,¹ Lars R. Schreiber^{4,5},^{4,5}
Hendrik Bluhm^{4,5} and Wolfgang Lechner^{1,2,3}

¹Parity Quantum Computing Germany GmbH, Schauenburgerstraße 6, 20095 Hamburg, Germany

²Parity Quantum Computing GmbH, Rennweg 1, Top 314, 6020 Innsbruck, Austria

³Institute for Theoretical Physics, University of Innsbruck, 6020 Innsbruck, Austria

⁴JARA-FIT Institute for Quantum Information, Forschungszentrum Jülich GmbH and RWTH Aachen University, 52074 Aachen, Germany

⁵ARQUE Systems GmbH, 52074 Aachen, Germany



(Received 5 March 2024; revised 7 May 2024; accepted 12 July 2024; published 5 August 2024)

Motivated by the prospect of a two-dimensional square-lattice geometry for semiconductor spin qubits, we explore the realization of the parity architecture with quantum dots. We present sequences of spin shuttling and quantum gates that implement the parity quantum approximate optimization algorithm (QAOA) on a lattice constructed of identical unit cells, such that the circuit depth is always constant. We further develop a detailed error model for a hardware-specific analysis of the parity architecture, and we estimate the errors during one round of parity QAOA. The model includes a general description of the shuttling errors as a function of the probability distribution function of the valley splitting, which is the main limitation for the performance. We compare our approach to a superconducting transmon qubit chip, and we find that with high-fidelity spin shuttling the performance of the spin qubits is competitive or even exceeds the results of the transmons. Finally, we discuss the possibility of decoding the logical quantum state and of quantum error mitigation. We find that already with near-term spin qubit devices, a sufficiently low physical error probability can be expected to reliably perform parity QAOA with a short depth in a regime where the success probability compares favorably to standard QAOA.

DOI: [10.1103/PhysRevB.110.075302](https://doi.org/10.1103/PhysRevB.110.075302)

I. INTRODUCTION

In the pursuit of quantum computing, the question of the ideal hardware platform is still unanswered. Among the contestants, spin qubits in gate-defined quantum dots (QDs) [1] stand out with their unique promise of leveraging the sophisticated manufacturing capabilities of the semiconductor industry once a design based on scalable building blocks is devised [2–4]. This, combined with their small size of only a few tens of nanometers per QD, may allow for the fabrication of quantum computers with millions of qubits that can easily be mass-produced [5].

While the gate fidelities are approaching the threshold to quantum error correction [6–8], serious challenges must still be overcome on the road to spin-based quantum computing. The most prominent are environmental electric noise [9,10], cross-talk and residual exchange interaction within dense arrays of qubits [11–14], and the demanding space requirements of the voltage gates and control electronics [2–4].

A possible means of addressing those problems could be spin shuttling, coherently moving the qubits between sites—with different functionality—on the chip on demand. Shuttling can be realized either in the conveyor-mode where a sliding potential well smoothly displaces the qubit [15,16], or as a bucket brigade by coherent tunneling between adjacent QDs [17–19]. While the latter variant requires a high degree of individual control [20,21], conveyor-mode shuttling with potentials formed by dedicated gates is showing promise for success [16,22,23]. In silicon heterostructures, the degenerate

conduction-band minima lead to an additional pseudospin, namely the valley degree of freedom, whose splitting is determined by the microscopic properties of the interface [24,25]. Local minima of the valley splitting can be a major challenge for conveyor-mode shuttling, however their occurrence can be reduced by engineering the semiconductor heterostructure [26–31] or adjusting shuttling trajectories [32].

A major objective in the development of spin qubits is the creation of connectivity between qubits in two dimensions. Inevitably for error-corrected quantum computing [33], this milestone could also make the parity architecture [34,35] a viable way to advance the performance of spin qubits. In the parity architecture, the logical state of the quantum computer is encoded by physical qubits that represent the parity of the logical spins [see Figs. 1(a) and 1(b)]. While introducing a qubit overhead, this encoding removes the requirement for long-distance interactions, and the redundant information allows for quantum error mitigation [36] and quantum error correction [37] for bit-flip errors. Furthermore, it allows the execution of the parity quantum approximate optimization algorithm (QAOA) [38,39] and reduces the circuit depth of cornerstone algorithms such as the quantum Fourier transform [40,41].

The QAOA is a gate-based algorithm for solving combinatorial optimization problems on a digital quantum computer [42–44], inspired by adiabatic quantum computing, where a quantum state is evolved adiabatically under a Hamiltonian representing the cost function of the optimization problem in order to approximate its ground state [45]. Here,

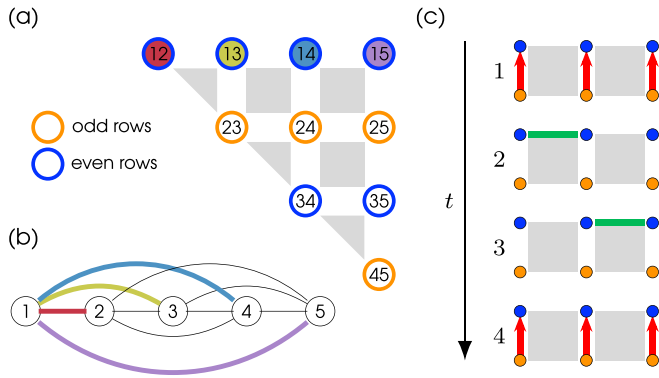


FIG. 1. (a) Visualization of the parity architecture. (b) An optimization problem on logical qubits 1–5 can be expressed by assigning one physical qubit ij in panel (a) to each logical interaction J_{ij} . This allows implementing the problem Hamiltonian with local fields, while constraints must be enforced on plaquettes of four or three adjacent qubits (gray). Note that it is also possible to compile higher-order interaction terms to an arbitrary chip layout [35]. For the decomposition of the constraints in parity QAOA, the qubits are organized in rows 0, 1, 2, ... (blue and orange). (c) Half of the circuit that enforces the constraints of the parity transformation during parity QAOA [39]. The lattice is separated into ribbons of adjacent rows on which CNOT (red arrows pointing from control to target) and ZZ gates (green lines) are applied in successive steps 1–4. All ribbons with an even-numbered top row can be treated in parallel, followed by all ribbons with an odd-numbered top row in parallel.

adiabatic time evolution is replaced by an alternating sequence of parametrized small-angle rotations corresponding to a problem and driver Hamiltonian, respectively. The parameters are then optimized in a quantum-classical feedback loop. The first proof-of-principle demonstrations of QAOA were shown on existing quantum hardware [46–48], and the algorithm may be a suitable candidate to prove a quantum advantage for nontrivial problems with a few hundred qubits and gate fidelity below the error correction threshold [43,49,50]. However, it remains challenging to achieve the connectivity required for an arbitrary problem Hamiltonian, which contributes to a nonoptimal scaling of resources with the problem size and complexity [47,50]. This problem is alleviated by the parity architecture [34,38].

We find that spin qubits can efficiently implement the parity architecture even compared to more mature platforms since their native gates naturally fit the demands for parity QAOA and thus require little additional transpilation. Different strategies are currently under investigation for realizing a two-dimensional lattice of spin qubits, which perfectly suits the parity architecture. Its possibility of working with local fields and nearest-neighbor interactions allows leveraging the fast, albeit short-range, two-qubit gates, one of the main advantages of spin qubits.

In this article, we investigate the performance of parity QAOA on two scalable spin qubit architectures based on shuttling and modularization of the chip. We develop the shuttling sequences for implementing the algorithm on these architectures, and we show that parity QAOA can be efficiently executed even though the topology of the chips is not a square lattice, as required in the original proposals. This result is of

relevance for many hardware platforms that suffer from a low connectivity, and it is of particular interest for spin qubits, where the realization of fully connected two-dimensional arrays is hindered by cross-talk and the space required for the fan-out of the voltage gates.

Generic error models have been used to analyze the performance of parity QAOA, although they do not allow us to reliably gauge the performance of real hardware. Here, instead, we introduce a realistic error model based on recent experimental, theoretical, and computational results for an in-depth performance analysis of parity QAOA. Based on our error model, the error probability of a single physical qubit is estimated. We find that with just slightly optimistic assumptions for the future development of the spin qubit coherence time, a full round of parity QAOA is feasible on both architectures, with the possibility to further enhance the performance by error mitigation. We note that the error probability is in a regime where simulations estimate a higher success rate for parity QAOA compared to standard QAOA [36], and where parity QAOA can address nontrivial problems [51].

Although the main result of the article—the analysis of the algorithmic performance—is specific to the hardware under consideration, this prediction of good performance of an algorithm tailored to the parity architecture under a realistic noise model suggests that the parity architecture is a suitable candidate for demonstrations on different types of noisy near-term quantum hardware. In particular, we make a comparison with a chip consisting of capacitively coupled transmons in the same layout as the modular spin qubit architecture. For an optimal shuttling velocity and an engineered valley splitting, the performance of parity QAOA with spin and superconducting qubits is comparable, and a shuttling-based spin qubit processor can surpass even optimistic assumptions for the transmon qubits.

The remainder of this article is organized as follows. In Sec. II a brief introduction to parity quantum computing is provided, followed by our proposed implementation of the parity QAOA algorithm on a spin qubit quantum processor in Sec. III. In particular, Sec. III A focuses on the case of a sparse shuttling-based architecture, and Sec. III B focuses on the case of a modular architecture where the registers are connected by spin shuttling. In Sec. IV the error model for all relevant processes is introduced. In Sec. V A the performance of the QAOA on both architectures is investigated in the presence of realistic errors, and in Sec. V B the results are put into context with superconducting transmon qubits. In Sec. VI the possibilities for error mitigation and the decoding of the quantum state are discussed. Finally, in Sec. VII the results are summarized and the paper is concluded.

II. INTRODUCTION TO PARITY QUANTUM COMPUTING

Here we provide an introduction to the parity architecture in general (Sec. II A) and the parity QAOA in particular (Sec. II B).

A. Parity architecture

Consider an optimization problem encoded in a spin-glass problem Hamiltonian with N qubits and K interactions of the

form

$$H_{\text{op}} = \sum_{i=1}^N J_i \sigma_z^{(i)} + \sum_{j<i} J_{ij} \sigma_z^{(i)} \sigma_z^{(j)} + \sum_{k<j<i} J_{ijk} \sigma_z^{(i)} \sigma_z^{(j)} \sigma_z^{(k)} + \dots, \quad (1)$$

where the coefficients $J_i, J_{ij}, \dots$ denote the coupling strengths, and $\sigma_z^{(i)}$ denotes the Pauli-Z operator acting on logical qubit i . The parity architecture [34,35] maps this Hamiltonian to the parity Hamiltonian

$$H_{\text{parity}} = \sum_{m=1}^K \tilde{J}_m \tilde{\sigma}_z^{(m)} + c \sum_{l=1}^{K-N+D} C_l \quad (2)$$

with K physical qubits, where each interaction between logical qubits in H_{op} was mapped to a local field term with strength $\tilde{J}_m$ on physical qubit m with associated Pauli-Z operator $\tilde{\sigma}_z^{(m)}$, e.g., $J_{ij} \sigma_z^{(i)} \sigma_z^{(j)} \mapsto J_m \tilde{\sigma}_z^{(m)}$, and $c > 0$ denotes a constant. The physical qubits represent the parity of a set of logical qubits, i.e., their σ_z -eigenvalue indicates if the respective logical qubits are in the same or opposite eigenstate. As the Hilbert space is thereby enlarged and contains nonphysical states, $K - N + D$ constraints

$$C_l = -\tilde{\sigma}_z^{(l_1)} \tilde{\sigma}_z^{(l_2)} \tilde{\sigma}_z^{(l_3)} [\tilde{\sigma}_z^{(l_4)}] \quad (3)$$

on three or four qubits, represented by the last sum in H_{parity} , are introduced [34,35]. Here, D denotes the number of ground-state degeneracies of H_{op} , and the square brackets in Eq. (3) indicate that the fourth qubit involved in the interaction is optional. This mapping is illustrated in Fig. 1(a).

The indices l_i in Eq. (3) are chosen such that all the logical indices involved in C_l occur an even number of times across all l_i . The constraints C_l stabilize the space of logical states, i.e., of states that have a correspondence in H_{op} and are therefore valid. Crucially, the interactions C_l can be implemented between qubits in geometrical proximity on a two-dimensional (2D) grid with nearest-neighbor connectivity [see Fig. 1(a)]. Therefore, only geometrically local interactions are required for solving optimization problems of arbitrary order on digital or analog quantum devices, which is particularly important in the noisy intermediate-scale quantum (NISQ) era [49] where implementing long-range interactions remains a challenge. Furthermore, the implementation of the constraint operators on digital quantum computers can be parallelized very efficiently. It was shown that the parity architecture is suitable for analog [34] and digital quantum optimization [38] as well as universal quantum computing [37]. We note that the parity architecture can also be viewed as an error correction code for bit-flip errors, at the cost of implementing nontransversal logical R_x rotations, resulting in nonlocal physical gates. However, for the NISQ algorithm, such as the QAOA, we do not strive for full error correction, which is why we can exploit the parity architecture to *remove* long-range interactions. Therefore, performing the parity transformation and adding the constraint operators to the resulting Hamiltonian allows for an implementation of quantum annealing and the quantum approximate optimization algorithm [42] with only local operations, and, for the latter, in constant circuit depth.

B. Parity QAOA

The quantum approximate optimization algorithm (QAOA) [42] aims at solving optimization problems encoded in an N -qubit problem Hamiltonian H_{op} by preparing a solution candidate state

$$|\psi(\boldsymbol{\beta}, \boldsymbol{\gamma})\rangle = \prod_{j=1}^p U_x(\beta_j) U_{\text{op}}(\gamma_j) |+\rangle \quad (4)$$

for an optimization problem on N qubits. Here, the initial state $|+\rangle$ denotes the state where all qubits i are in an equal superposition of the eigenstates of σ_z , $(|0\rangle_i + |1\rangle_i)/\sqrt{2}$, the unitary

$$U_{\text{op}}(\gamma) = e^{-i\gamma H_{\text{op}}} \quad (5)$$

represents the time evolution operator under the problem Hamiltonian, and

$$U_x(\beta) = \prod_i e^{-i\beta \tilde{\sigma}_x^{(i)}} \quad (6)$$

is the so-called driver unitary. The $2p$ parameters $\boldsymbol{\beta} = (\beta_1, \dots, \beta_p)$ and $\boldsymbol{\gamma} = (\gamma_1, \dots, \gamma_p)$ are optimized in a quantum-classical feedback loop by using a quantum computer to compute the energy expectation value $\langle H_{\text{op}} \rangle = \langle \psi(\boldsymbol{\beta}, \boldsymbol{\gamma}) | H_{\text{op}} | \psi(\boldsymbol{\beta}, \boldsymbol{\gamma}) \rangle$, and employing a classical routine to optimize $\boldsymbol{\beta}$ and $\boldsymbol{\gamma}$ with respect to $\langle H_{\text{op}} \rangle$ until some termination criterion is reached. The candidate ground state obtained as an optimization result is then determined by reading out all qubits. By using the parity Hamiltonian H_{parity} as the problem Hamiltonian and separating the local field term and the constraint term into two separate QAOA unitaries, extending the search space by $\boldsymbol{\Omega} = (\Omega_1, \dots, \Omega_p)$ to $3p$ classical parameters, parity QAOA [38] is obtained.

Note that it is possible to implement the QAOA in the parity architecture without the constraint terms by allowing more complex and nonlocal driver operators [52]. However, this approach exhibits a circuit depth growing linearly with the system size. In this work, we stick to the original proposal of explicitly implementing the C_l terms in this work, and we exploit the geometric locality of these operators.

In its explicit form, parity QAOA requires single-qubit operations on all qubits and an implementation of the three- and four-body constraints arising from the parity transformation, thus preparing the final state [38]

$$|\psi(\boldsymbol{\beta}, \boldsymbol{\gamma}, \boldsymbol{\Omega})\rangle = \prod_{j=1}^p U_z(\gamma_j) U_c(\Omega_j) U_x(\beta_j) |+\rangle, \quad (7)$$

where the unitary

$$U_z(\gamma) = e^{-i\gamma \sum_i \tilde{J}_i \tilde{\sigma}_z^{(i)}} \quad (8)$$

is the time evolution under the problem Hamiltonian after the parity transformation. The operator

$$U_c(\Omega) = e^{-i\Omega \sum_l \tilde{\sigma}_z^{(l_1)} \tilde{\sigma}_z^{(l_2)} \tilde{\sigma}_z^{(l_3)} \tilde{\sigma}_z^{(l_4)}} \quad (9)$$

enforces the constraints on all plaquettes l . This total time evolution is repeated for a number of rounds p .

The operators U_x and U_z can be trivially implemented with single-qubit rotations, and it has been shown [38,39] that also U_c can be implemented in constant depth. We choose the

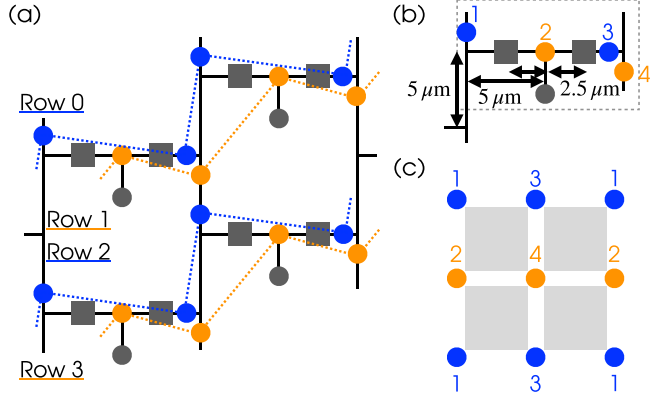


FIG. 2. The spin bus platform under consideration [4]. (a) The electrons (blue and orange dots) are stored in shuttling lanes (black lines); qubits are labeled 1–4 in each unit cell. The electrons can be conveyed to manipulation (dark gray squares) or initialization/readout zones (dark gray circles); T-junctions provide connectivity between qubits in two dimensions. The vertical shuttling lanes are controlled globally. The segment holding qubits 1–4 can be considered a unit cell of the lattice. For the purpose of the QAOA, the electrons are separated in even (blue) and odd (orange) rows indicated by the dashed lines. They correspond to the rows from Fig. 1. (b) One unit cell with four qubits, highlighted by the dashed box. We assume that qubits 1, 3, and 4 are placed $1.25\,\mu\text{m}$ from the T-junction (and the manipulation zone, in the case of qubit 3). The classical control electronics and the gate fan-out can find space between the unit cells. (c) Corresponding plaquettes of four-qubit constraints (light gray) of the parity architecture. The labels denote the qubits' position in their unit cell.

decomposition from Ref. [39], which is shown in Fig. 1(c), to achieve the latter. This approach separates the square grid of qubits into ribbons of two neighboring rows each. The constraints are then enforced by a sequence of CNOT and ZZ gates, which can be executed in parallel on each ribbon and furthermore allows parallelization of all ribbons that include an even-numbered top row of qubits in the first step and ribbons with odd top rows in the second steps. The rows are indicated in Figs. 1, 2, and 5. In this work, for simplicity, we consider the case with an implementation of all constraints in the sequence of eight time steps from Fig. 1(c), where at each boundary between two plaquettes both qubits are included in the respective constraints [34] (i.e., square plaquettes only). More general problems [35] can be treated by including additional CNOT gates between steps 2 and 3 [39].

Apart from the advantages of solely local interactions and a constant circuit depth for QAOA, the parity architecture offers the advantage of an intrinsic possibility for error mitigation. The redundant information held by the additional qubits can be exploited to read out several spanning trees of the logical system in order to detect and correct constraint violations, either due to quantum errors or because of constraint violations by the mixer unitary [36]. The latter can also be avoided by using constraint-preserving driver operators [52], however at the cost of giving up the constant circuit depth. Even though physical errors reduce the success rate for finding the ground state, it was shown that the decoding of the spanning trees

still grants a high success rate and can outperform standard QAOA [36].

III. PARITY QAOA WITH SPIN QUBITS

In this section, we present an implementation of the parity QAOA on two possible two-dimensional electron spin qubit architectures: an extremely sparse architecture entirely based on spin shuttling [4] in Sec. III A, and an architecture with small, dense registers of QDs connected by shuttling links [2] in Sec. III B. Here, the mapping to the chip layout and the compiled shuttling and gate sequences are given; a comparison of the performances will follow in Sec. V A. Animated versions of the gate and shuttling sequence can be found in the supplemental material [53].

A. Implementation on a sparse spin bus architecture

Recently, the so-called spin bus architecture, Fig. 2, was proposed [4]. There, the electrons are stored in shuttling lanes defined by periodically interconnected voltage gates, which allow a smooth conveyor-mode shuttling of the charge by applying phase-shifted sinusoidal voltages to each set of connected gates [16,22,23]. The shuttling lanes form a two-dimensional lattice with dedicated manipulation and initialization/readout zones to which the electrons can be shuttled on demand. This architecture has the advantage of high connectivity and promises a long coherence time, since the electrons can be stored far from the detrimental effects of the micromagnets at the manipulation zones. Due to its sparse nature, it can be expected to suffer from little crosstalk and it creates space for the voltage gates and classical control electronics. The architecture is sketched in Fig. 2.

Unlike Ref. [4], we assume that each unit cell contains four electrons and two manipulation zones, doubling the number of qubits at the cost of the additional input lines for one conveyor per cell. Thus, the number of qubits is doubled with only a moderate increase in complexity. We expect this to be beneficial for near-term devices limited by the requirements of room-temperature control. Furthermore, the increased density of qubits results in shorter shuttling paths within the unit cells, thus mitigating shuttling-related errors. Alternatively, it would also be possible to carry out the operations sequentially while storing electrons not involved in shuttling lanes, thus reducing complexity at the cost of a slightly reduced performance.

The constraints are implemented by the gate sequence shown in Fig. 3 for the set of ribbons consisting of even top and odd bottom rows, followed by the circuit in Fig. 4 for the set of ribbons consisting of odd top and even bottom rows. Due to the nonequivalent positions of the qubits in the unit cell, these two cases need to be treated separately. The shuttling and gate operations can be performed in parallel on all unit cells, achieving an optimal circuit depth. Due to the global control of the vertical shuttling lanes, some electrons are shuttled unnecessarily, as shown in Fig. 3. However, with high-fidelity shuttling sufficient for several tens of micrometers, we do not expect this short excess distance to be problematic; averaging quasistatic noise by shuttling (motional narrowing) might even enhance the coherence of an idle qubit.

During the constraint circuit, every bulk qubit—which is not located on the edge of the physical device—is a target

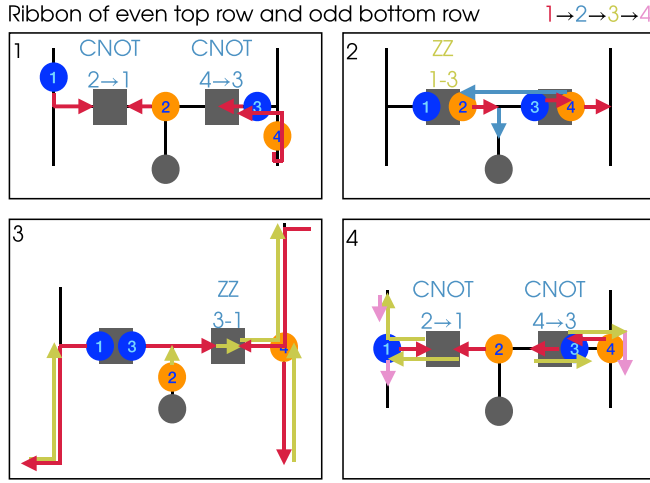


FIG. 3. Circuit for implementing the parity constraints on ribbons of an even top and an odd bottom row if executed on each unit cell in parallel. The arrows indicate the operation of spin conveyors; gate operations are indicated near the manipulation zone. The coloring represents the time ordering of the sequence: from red to light blue to lime to pink as indicated in the top right corner. Due to the global vertical shuttling lanes, some electrons are moved unnecessarily, for example qubit 4 in steps 1 and 3 and qubit 1 in step 4. Note that qubit 1 leaves its unit cell and undergoes a gate in a different cell in step 3, then returns in step 4, such that the second ZZ gate is between qubits from different unit cells. Steps 1–4 correspond to steps 1–4 in Fig. 1(c).

and control of two CNOT gates each and participates in two ZZ gates, while qubits at the boundary of the lattice experience fewer gates. Qubit 1 is shuttled over a distance of $42.5 \mu\text{m}$ with our assumptions on the size of the unit cell and idles for $46.25 \mu\text{m}/v + 2T_{ZZ}$, where v is the shuttling velocity and T_{ZZ} is the time required for a ZZ gate. Qubit 2 is shuttled over $62.5 \mu\text{m}$ and idles for $26.25 \mu\text{m}/v + 2T_{ZZ}$, qubit 3 is shuttled over $26.25 \mu\text{m}$ and idles for $62.5 \mu\text{m}/v + 2T_{ZZ}$, while qubit 4 moves by $61.25 \mu\text{m}$ and idles for $27.5 \mu\text{m}/v + 2T_{ZZ}$. This is not counting the shuttling of step 1 for the even-odd ribbons since it can be absorbed into the single-qubit gates of the driver term. The shuttling velocity v is a critical parameter for the duration of the algorithm and also for the strength of the errors, as will be discussed in Sec. IV. The discrepancy in shuttled distance between the qubits is due to the fact that qubits 2 and 4 are moved to a foreign unit cell in order to implement the constraints on the odd-even ribbons, while qubit 3 is favorably placed close to two manipulation zones and is thus highly connected without much movement.

The circuit for the execution of single-qubit gates, required for implementing U_x and U_z , is shown in Fig. 4 in the panel SQG. In a single-qubit gate step followed by the constraint circuit, qubits 1 and 3 can be returned to the manipulation zone after the operation on qubits 2 and 4 is complete, thus simplifying step 1 from the even-odd rows, Fig. 3. Alternatively, all qubits can be returned to their idling position, e.g., to wait until they are read out at the end of the algorithm. The former version adds $8.75 \mu\text{m}$ ($2.5 \mu\text{m}$, $6.25 \mu\text{m}$, $6.25 \mu\text{m}$) of shuttling on qubit 1 (2, 3, 4), while the latter adds $10 \mu\text{m}$ ($5 \mu\text{m}$, $7.5 \mu\text{m}$, $10 \mu\text{m}$) on qubit 1 (2, 3, 4).

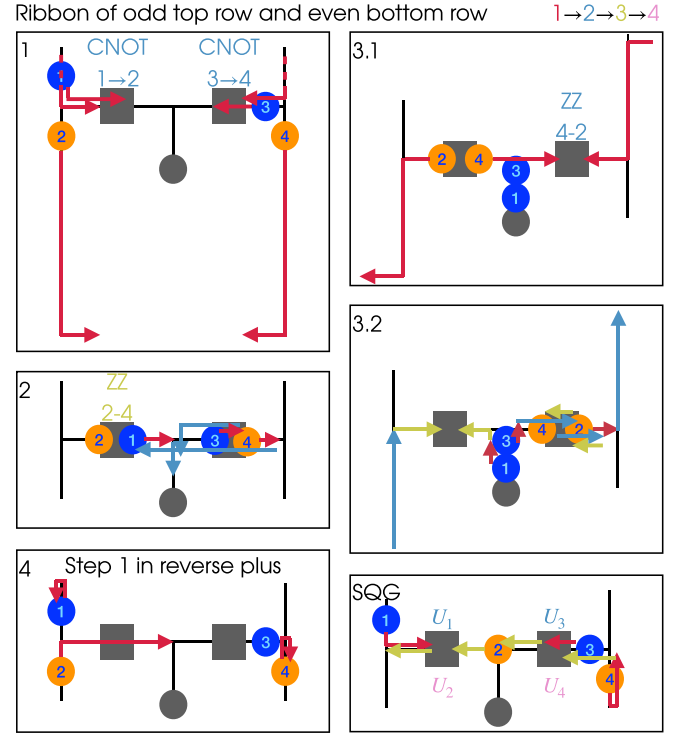


FIG. 4. Circuit for implementing the parity constraints on ribbons of an odd top and an even bottom row (1–4) if executed on each unit cell in parallel, as well as for single-qubit gates (SQGs). To implement the constraints on this set of ribbons, qubits 2 and 4 are shuttled to an adjacent unit cell in step 1 and return in step 4. The third step is separated in 3.1, which includes the gate, and step 3.2, which restores the final configuration of step 1 such that the final CNOT gate and the shuttling in step 4 can be performed easily. The circuit SQG for implementing arbitrary single-qubit gates U_i on all qubits i can be followed by returning the qubits to their home position (initial configuration of step 1 in Fig. 3), or by returning qubits 1 and 3 to the manipulation zone (final configuration of step 1 in even-odd ribbons) if the single-qubit gate is followed by a constraint step.

Initialization and readout, the remaining building blocks of any algorithm, are performed by successively shuttling each qubit from or to the initialization/readout zone. To initialize each qubit to an arbitrary state, electrons 1, 3, and 4 can be stopped at the manipulation zone they pass on the way to their starting position, and qubit 2 can be routed on a detour to either manipulation zone.

B. Implementation on a modular architecture

A tradeoff between the typical dense array of QDs [54] and the extremely sparse spin bus is an architecture where dense registers of a few qubits are connected by coherent quantum links, as depicted in Fig. 5 [2,3]. This retains the advantages of dense arrays—fast gates and a small footprint—while creating space for the fan-out of voltage gates and classical electronics. Here, we consider a minimal version of this modular architecture, where registers of 2×2 QDs are connected via diagonal conveyor-mode shuttling lanes in two dimensions, as depicted in Fig. 5.

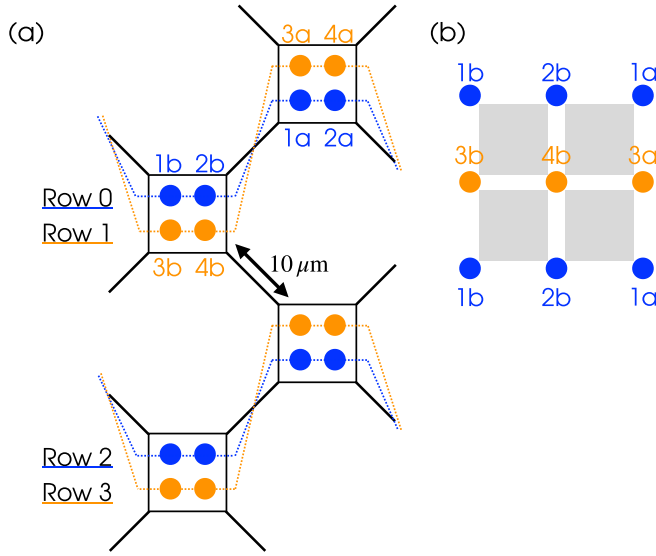


FIG. 5. The modular architecture under consideration [2]. (a) Registers of 2×2 QDs are interconnected via diagonal shuttling lanes. Manipulations are performed on-site, and we assume that one site per register is equipped for readout, shuttling serves for connectivity only. Here, a unit cell consists of two registers denoted by a and b with eight qubits in total, since it is favorable to mirror the registers alternately. This architecture provides space for classical electronics and the gate fan-out in between the registers. (b) Corresponding plaquettes of the parity architecture. The labels denote the qubits' position in their unit cell.

A sequence of operations for implementing the constraint term on this architecture is shown in Figs. 6 and 7. Here, a number of SWAP operations cannot be avoided since each qubit is coupled to only three direct neighbors, which is not ideal for realizing the parity architecture. Again, we find that for the constraint term all bulk qubits are controls and targets of two CNOT gates each, and they participate in two ZZ gates and, additionally, ten SWAP gates. Furthermore, all qubits at positions 1 and 3 in the unit cells are shuttled by $40 \mu\text{m}$ and idle

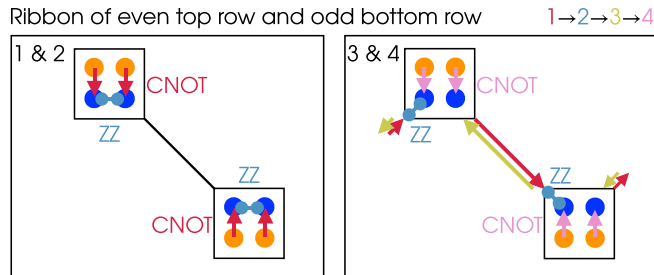


FIG. 6. Circuit for implementing the parity constraints on ribbons of even top and odd bottom rows in a modular architecture if executed on each unit cell in parallel. Arrows and lines within the registers indicate two-qubit gates, while arrows between the registers indicate the operation of shuttling lanes. The time ordering is indicated by the colors, from red to light blue to lime to pink as indicated in the top right corner. The steps 1–4 here correspond to the decomposition of the circuit in Fig. 1(b) and on the spin bus in Fig. 3.

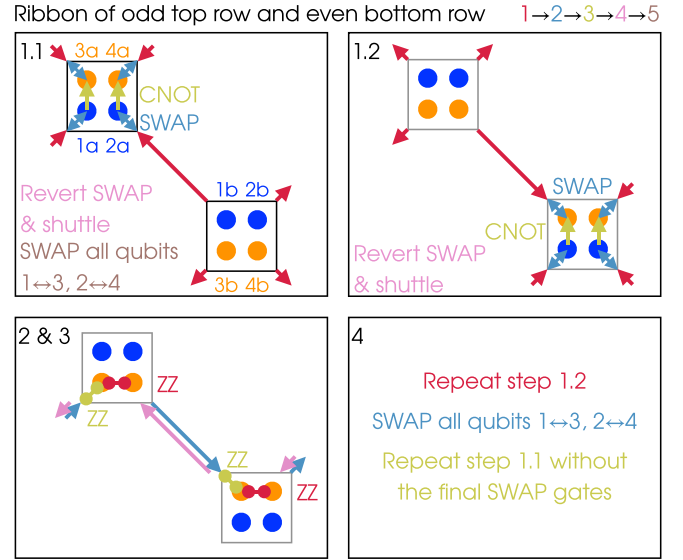


FIG. 7. Circuit for implementing the parity constraints on ribbons of an odd top and an even bottom row in a modular architecture if executed on each unit cell in parallel. Here, the lower connectivity compared to the spin bus is apparent, as each qubit is subject to a total of 10 SWAP gates in order to implement the required interaction. The steps 1–4 correspond to those in Fig. 1(b) and the implementation on the spin bus in Fig. 4. In steps 1 and 4, the pulses for the gates can be applied globally to all registers since all QDs in the nonaddressed register are vacant. Note that in steps 1 and 4 the qubits $1a$ – $4a$ and $1b$ – $4b$ in the mirrored registers undergo the same operations, albeit in different order.

for $80 \mu\text{m}/v + 2T_{ZZ} + 2T_{\text{CNOT}}$, while all qubits at positions 2 and 4 are shuttled by $60 \mu\text{m}$ and idle for $60 \mu\text{m}/v + 2T_{ZZ} + 2T_{\text{CNOT}}$. It is possible to replace eight SWAP gates per qubit with hopping between neighboring QDs within the registers if an alternative circuit is used and empty registers are available at the edges of the lattice (see Appendix A).

Single-qubit gates can be performed on-site in each QD and do not require shuttling, and we assume that one QD per register is equipped with a readout apparatus, allowing us to measure and initialize all qubits in four successive steps combined with SWAP gates or other operations to transfer the spin projection between neighboring QDs. In summary, time and operations are comparable on both architectures, with the addition of 10 SWAP gates in the modular architecture. The spin bus is able to implement a single round of QAOA $\approx 10 \mu\text{m}/v$ faster than the modular architecture, including initialization and readout.

IV. GATE AND ERROR MODEL

In this section we introduce our error model, which is applied to both hardware layouts. There are three general sources of errors:

- (i) Errors from the quantum operations on the qubits.
- (ii) Errors from the shuttling processes.
- (iii) Errors from initialization and readout.

For the specific details of the hardware platform, we assume electron spins confined in QDs in a Si/SiGe heterostructure. This choice is motivated by the outstanding

coherence properties of QDs in isotopically purified silicon [1], the demonstration of high-fidelity quantum gates [6–8], and the recent successes of electron conveyors in this material [22,23]. Thus, all necessary building blocks for our architecture are available.

A. Gate errors

We assume that each manipulation zone in the spin bus architecture and each register in the modular architecture is equipped with a micromagnet. Thus, single spin manipulation can be accomplished by means of electric dipole spin resonance. Lowering the tunnel barrier between two adjacent QDs gives rise to a nearest-neighbor exchange interaction [1]. We decompose all two-qubit gates into gates that are well-proven and optimized [1,6,7].

The gate errors are modeled using the Kraus representation of imperfect quantum channels [55] where the error probabilities are chosen in order to describe realistic error rates. For single-qubit gates, we assume a depolarization channel with a probability of $p_d = 10^{-3}$ [4,6–8], thus the density operator ρ_f after a single-qubit gate U_i on qubit i is given by $\rho_f = \sum_k K_{k,i} \rho_0 K_{k,i}^\dagger$, where $K_{0,i} = \sqrt{1 - p_d} U_i$ and $K_{k,i} = \sqrt{p_d/3} \tilde{\sigma}_k^{(i)}$ for $k = x, y, z$ with the Pauli operators $\tilde{\sigma}_k^{(i)}$ of qubit i and the input density matrix ρ_0 .

The default entangling two-qubit gate in QDs equipped with a micromagnet is the controlled phase gate, $\text{CP}_\alpha = \text{diag}(1, 1, 1, e^{i\alpha})$, in the basis $\{|0, 0\rangle, |0, 1\rangle, |1, 0\rangle, |1, 1\rangle\}$ [6,7]. Similar to the single-qubit gates, we assume that the errors of the CP gate are captured by a phase-flip channel with probability $p_\phi = 10^{-3}$ and a bit-flip channel with probability $p_b = 10^{-6}$ on both qubits, taking into account that dephasing is the limiting error mechanism for spin qubits. However, the circuits presented in the previous section require the two-qubit gates ZZ, CNOT, and SWAP. These are synthesized from CP and single-qubit gates by concatenating the respective quantum channels.

The gate $\text{ZZ}_\alpha^{(i,j)} = \exp(i\alpha \tilde{\sigma}_z^{(i)} \tilde{\sigma}_z^{(j)})$ between qubits i and j with rotation angle α is obtained from the decomposition

$$\text{ZZ}_\alpha^{(i,j)} = \text{CP}_{-2\alpha}^{(i,j)} R_z^{(i)}(\alpha) R_z^{(j)}(\alpha) \quad (10)$$

with the single-qubit rotation $R_z^{(i)}(\alpha) = e^{-i\alpha \tilde{\sigma}_z^{(i)}/2}$ around the z axis of qubit i . Note that the ZZ gate and its error are symmetric between the two qubits.

Analogously, the CNOT gate with control qubit i and target j can be represented as

$$\text{CNOT}^{(i,j)} = H^{(j)} \text{CP}_\pi^{(i,j)} H^{(j)} \quad (11)$$

with the Hadamard gate H . A SWAP gate is then obtained by a sequence of three CNOTs:

$$\text{SWAP}^{(i,j)} = \text{CNOT}^{(j,i)} \text{CNOT}^{(i,j)} \text{CNOT}^{(j,i)}. \quad (12)$$

These composite gates can reliably be performed with a high fidelity in Si/SiGe quantum dots with micromagnets [6,7], although we expect that it is possible to engineer a more efficient version of the SWAP gate realized by the exchange coupling or by including a physical position swap via shuttling.

Qubits idling for a time t will suffer from dephasing with a characteristic timescale T_2 and relaxation with a characteristic timescale T_1 due to environmental noise. A conservative lower bound for T_2 is the pure dephasing time T_2^* . These are modeled with the Kraus representation of a phase damping and amplitude damping channel [55]. For $t \ll T_2, T_1$ the probabilities for these channels can be approximated as $p_{\phi,\text{idl}} = t/T_2$ ($p_{r,\text{idl}} = t/T_1$) for dephasing (relaxation). If the dephasing is dominated by quasistatic noise, the decay of the coherences is described by a Gaussian by the choice $p_{\phi,\text{idl}} = (t/T_2)^2$ [1].

As an optimistic estimate, we use $T_1 = 1$ s for both architectures and $T_2 = 20$ μs for the modular layout [56], although this number can improve significantly if dynamical decoupling is applied [1]. The spin bus allows storing the electrons further away from the detrimental field of the micromagnet, which couples the spin to electric field fluctuations, as well as the SETs and charge reservoirs, which are sources of Johnson noise. Thus a longer dephasing time can be expected. Using SiMOS QDs without a micromagnet as a reference, we take $T_2 = 100$ μs as an optimistic estimate for the spin bus layout [1]. Note that in both architectures, shuttling has an effect similar to dynamical decoupling and can increase a qubit's coherence time due to motional narrowing [21,22]. This is particularly relevant if a qubit is shuttled back and forth to and from a manipulation zone along the same path: Inverting the qubit state before the return allows the removal of certain adiabatic effects, such as deterministic variations of the qubit frequency during the shuttling [21]. This is trivially done in the modular architecture and is possible for most shuttling paths in the spin bus. Idling mostly occurs while qubits are waiting for the completion of shuttling processes or gates on other qubits. Typical timescales are $T_{1q} = 100$ ns for the duration of a single-qubit gate and $T_{2q} = 50$ ns for the duration of a native two-qubit gate [4].

B. Shuttling errors

In both architectures we assume conveyor-mode shuttling of electrons, and we have to take into account that the chosen host material exhibits near-degenerate valleys [24,25,57]. A detailed derivation of shuttling errors can be found in Ref. [21]; however, we extend the model for the effects of the valley pseudospin in order to take into account the recent discoveries concerning the dependence of the valley splitting on the alloy disorder [27–29].

Both valley splitting E_v and valley phase φ will follow a random trajectory along the shuttling path. Assuming that spin-valley relaxation hot spots [58–60] are avoided, the main effect of the valley is dephasing [21]. Due to nonadiabatic shuttling, the electron may have a random weight in the excited valley state and thus accumulate an unpredictable phase. To capture this effect, we assign each location x along the shuttling path a valley Hamiltonian

$$H'_v(x) = \frac{E_v(x)}{2} (e^{i\varphi(x)} |v_- \rangle \langle v_+| + e^{-i\varphi(x)} |v_+ \rangle \langle v_-|), \quad (13)$$

with the $\pm z$ valley states $|v_\pm \rangle$ [25]. At each point x the valley splitting and phase are drawn from distributions P_{E_v} and P_φ , respectively.

Since the position is time-dependent, $x = vt$ with shuttling velocity v , transforming H'_v into its instantaneous eigenbasis with the unitary U leads to a term that causes transitions between the instantaneous valley eigenstates,

$$H_v = U^\dagger H'_v U + i\hbar \dot{U}^\dagger U \quad (14)$$

$$= \frac{E_v}{2} \tau_z + \frac{\hbar \dot{\varphi}}{2} (\mathbb{1} - \tau_x), \quad (15)$$

where $\tau_z = |e_v\rangle\langle e_v| - |g_v\rangle\langle g_v|$ and $\tau_x = |g_v\rangle\langle e_v| + |e_v\rangle\langle g_v|$ are the Pauli z and x matrices for the instantaneous valley eigenstates. Approximating the differentiation as a quotient results in $\dot{\varphi}(t) = \Delta\varphi/\Delta t = v\Delta\varphi/\Delta x$, where Δx is the minimal distance at which a different valley splitting is resolved, and $\Delta\varphi$ is the difference in valley phase over this distance. A reasonable assumption for Δx is the dot size [32]. The difference $\Delta\varphi$ is a random variable whose probability distribution function can be obtained from a convolution of the distribution functions of the summands,

$$P_{\Delta\varphi} = \int_{-\pi}^{\pi} d\phi P_\varphi(\phi) P_\varphi[-(\Delta\varphi - \phi)]. \quad (16)$$

In the limit of a large valley splitting relative to the variation of the valley phase, $E_v/\hbar\dot{\varphi} \gg 1$, the probability of finding the electron in its excited valley state after moving it over a distance of Δx is easily obtained from the solution of the Schrödinger equation for a two-level system in the adiabatic frame [21],

$$p_{e,v} = |\langle e_v | e^{-iH_v \Delta x / \hbar v} | g_v \rangle|^2 \quad (17)$$

$$= \frac{(\hbar v \Delta\varphi / \Delta x)^2}{E_v^2 + (\hbar v \Delta\varphi / \Delta x)^2} \sin^2(\theta), \quad (18)$$

$$\theta = \sqrt{E_v^2 + (\hbar v \Delta\varphi / \Delta x)^2} \frac{\Delta x}{2\hbar v}. \quad (19)$$

Given the probability distributions, the average excitation probability over the distance Δx is thus

$$\bar{p}_{e,v} = \int_0^\infty dE_v \int_{-\pi}^{\pi} d(\Delta\varphi) P_{E_v}(E_v) P_{\Delta\varphi}(\Delta\varphi) p_{e,v}. \quad (20)$$

Consequently, the average probability for finding the electron in the excited valley state after moving a distance $L = n\Delta x$ is given by

$$p_v = 1 - (1 - \bar{p}_{e,v})^n \approx \bar{p}_{e,v} L / \Delta x. \quad (21)$$

In Ref. [29], a Rice distribution was found for the valley splitting E_v , and the numerical results of the same reference suggest $P_\varphi \approx 1/2\pi$ and thus $P_{\Delta\varphi} \approx 1/2\pi$. Consequently, p_v only depends on the two parameters of the Rice distribution—corresponding to the mean valley splitting and its variance—the shuttling velocity, dot size, and traveled distance. Our model for p_v is easily adapted to more accurate distribution functions unveiled by future research and hardware-specific distributions measured from individual devices by simply evaluating Eq. (20).

In accordance with Ref. [21] we then include an adiabatic contribution

$$p_{\text{ad}} = 2l_c^{\delta\omega} L / (vT_2)^2 \quad (22)$$

to the dephasing, due to fluctuations of the spin splitting with the correlation length $l_c^{\delta\omega}$. The fluctuations of the spin

splitting can originate from surrounding nuclear spins, magnetic field gradients, and spin-orbit coupling. This expression accounts for motional narrowing, which partly protects the shuttled spin. However, this effect is expected to be reduced in isotopically purified silicon, since the electric fluctuations dominating the noise in the absence of nuclear spins have a comparably large correlation length. Note that the effect of deterministic and reproducible variations of the spin splitting can be removed by calibration.

We choose a dot size of $\Delta x = 20$ nm and estimate a noise correlation length of $l_c^{\delta\omega} = 1$ μm . The latter is very conservative, but the effect of motional narrowing is strongly suppressed in that order of magnitude already. The dephasing during the shuttling is then described by the Kraus representation of a dephasing channel with the probability

$$p_{\text{deph}} = 1 - (1 - p_v)(1 - p_{\text{ad}}) \approx p_v + p_{\text{ad}}. \quad (23)$$

Relaxation of a shuttled spin is described by an amplitude damping channel with the probability $L/vT_1 + 10^{-4}L/10$ μm , where the second term emerges due to spatially varying transverse spin-orbit components [21].

Note that after averaging the valley effects, this is only the expected shuttling error in a typical shuttling lane. Due to its probabilistic nature, the valley splitting may strongly fluctuate between different shuttling lanes on a given device. Thus, all fidelity estimates obtained from this model represent an expectation value averaged over a large number of devices.

A number of possible shuttling errors are neglected in this description. These include the temporary breaking of a moving quantum well into a double well, which may harm the orbital state of the electron, the loss or capture of an electron, and the jumping of a moving electron to an adjacent empty well of the conveyor. This is justified if disorder in the device is sufficiently low. Experimental observations show that these types of errors are not detrimental to shuttling [16,22,23], and if relevant they can be reduced by technological optimization.

C. State preparation and measurement

In the spin bus architecture, readout and initialization are performed in dedicated zones. In the modular architecture, it is assumed that one QD per register is coupled to a readout/initialization apparatus. The starting point of the initialization is the singlet ground state of two electrons in an auxiliary QD, which is separated into two QDs by an adiabatic sweep. For readout, the same adiabatic sweep is performed in reverse, merging the electron to be measured into one dot with a reference electron. Due to the Pauli exclusion principle, only totally antisymmetric two-electron spin states are allowed in a single QD. Thus, depending on the regime of operation, this implements either a singlet-triplet measurement or a spin-parity measurement [1].

The adiabatic sweep can be performed with a fidelity as high as $F_m = 0.999$ [4], and the subsequent charge detection verifying the outcome of the sweep can be expected to have an error probability of $O(10^{-5})$ in devices optimized for spin shuttling [16]. The high fidelity of the adiabatic process requires a relatively long time, however. We assume a readout time of $T_r \approx 5$ μs per electron [4].

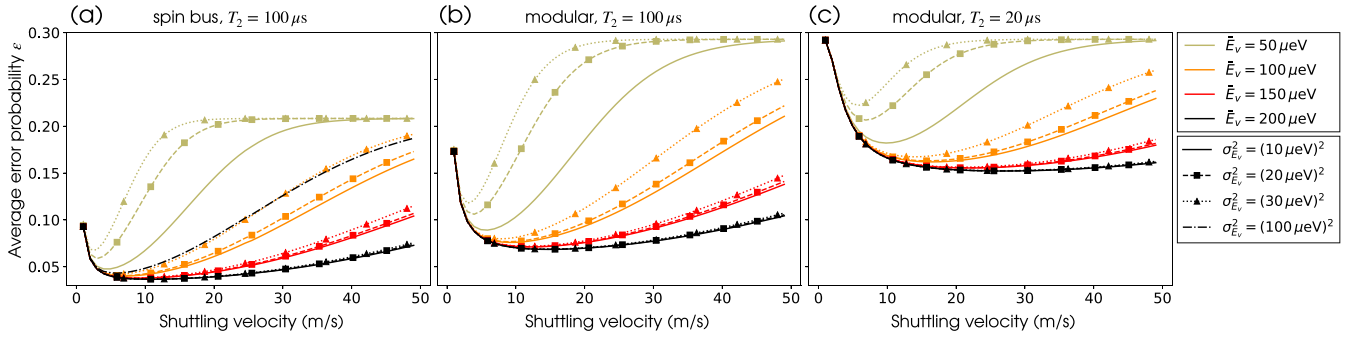


FIG. 8. Average error probability after one round of QAOA for (a) the spin bus with $T_2 = 100 \mu\text{s}$ and (b) the modular architecture with an optimistic $T_2 = 100 \mu\text{s}$, and (c) with a realistic $T_2 = 20 \mu\text{s}$. In all panels, the color (style) of each line encodes the mean $\bar{E}_v$ (variance $\sigma_{E_v}^2$) of the distribution of the valley splitting. The less weight the distribution has near $E_v = 0$, the better the algorithm performs and the broader the window of feasible shuttling velocities becomes. Since $T_2 \ll T_1$, the curves converge towards the scenario of a dephased qubit with almost intact spin projection.

V. PERFORMANCE ESTIMATES IN THE PRESENCE OF NOISE

In this section, we evaluate the gate sequences devised in Sec. III, including the errors described in Sec. IV. Subsequently, in Sec. VB, the result is put into context with other architectures, and a comparison with superconducting transmon qubits is given.

A. Performance of the spin qubit architectures

The time evolution of the qubit density matrix ρ is modeled by means of the Kraus operators for each operation. For simplicity, only one unit cell with periodic boundary conditions is modeled, undergoing one round of QAOA. This procedure will not return the correct output state of the algorithm, but is sufficient for estimating the physical errors by comparing the imperfect output state ρ_{err} with the ideal output state ρ_{id} without errors. Note that a single round of QAOA can deliver only a coarse approximation of the ground state, and generally the performance of QAOA improves with increasing number of rounds [46–48,50]. The fidelity

$$F = \left(\text{tr} \sqrt{\sqrt{\rho_{\text{id}}} \rho_{\text{err}} \sqrt{\rho_{\text{id}}}} \right)^2 \in [0, 1] \quad (24)$$

is a measure for the probability to find the unit cell in its desired output state, where no qubit suffered from an error [55]. Thus, we introduce the average single-qubit error probability p_{1q} . Assuming that errors on the four individual qubits are independent, p_{1q} is obtained from $F = (1 - p_{1q})^4$.

The single round of QAOA is followed by a readout of all qubits. The probability to observe an error on a single qubit at the end of the total circuit is thus

$$\varepsilon = 1 - (1 - p_{1q}) F_r F_m, \quad (25)$$

where F_r is the fidelity of the shuttling and idling in the readout step, and F_m is the fidelity of the measurement itself. This total physical single-qubit error probability is plotted for the spin bus architecture in Fig. 8(a) and for the modular architecture with both an optimistic assumption for T_2 that can

possibly be achieved by dynamical decoupling in Fig. 8(b) as well as with a realistic T_2 in Fig. 8(c).

The results here are obtained with $p_{\phi, \text{idl}} = (t/T_2)^2$. In the case of a Gaussian decay, $p_{\phi, \text{idl}} = (t/T_2)^2$ of the coherences due to low-frequency noise, the performance is found to be considerably better, in particular for slow shuttling. We discuss this case in Appendix B. Realistically, a result between those extremes can be expected.

Both architectures show the same general dependence on the distribution of the valley splitting and the shuttling velocity, although the spin bus performs slightly better even in case of identical dephasing [cf. Figs. 8(a) and 8(b)], despite the comparable amount of shuttling and idling. This is due to the fact that in the modular architecture a total of ten SWAP gates are required as additional steps, which are decomposed into 30 CP and 60 Hadamard gates, introducing additional errors, as discussed in Sec. III.

The error probability ε shows the interplay of the main dephasing mechanisms: If the algorithm requires a long time, the qubits will strongly dephase due to their finite T_2 . This is mitigated by increasing the shuttling velocity v , which lowers the error probability at first. However, as v increases, the nonadiabatic errors of the shuttling increase, such that all curves finally converge to the case of a fully decohered spin for large v . This leads to the emergence of an optimal shuttling velocity.

The probability distribution function of the valley splitting, characterized here by its mean $\bar{E}_v$ and variance $\sigma_{E_v}^2$, determines the strength of the nonadiabatic effects and thus the optimum. In particular, a distribution with a large weight near $E_v = 0$ is problematic for shuttling. Reducing the spread $\sigma_{E_v}^2$ of the distribution and increasing its mean $\bar{E}_v$ by engineering the interface of the semiconductor heterostructure will result in both a lower average error and a broader window of v in which near-optimal results can be expected. Increasing the mean valley splitting to $\bar{E}_v \approx 200 \mu\text{eV}$ with a standard deviation of $\sigma_{E_v} \approx 30 \mu\text{eV}$ or less—which is well within the theoretically predicted range of distributions [29]—a single-qubit error probability as low as $\varepsilon \approx 0.037$ ($\varepsilon \approx 0.069$) is observed for the spin bus (modular) architecture. If the noise is dominated by low-frequency components, even lower error probabilities are expected (see Appendix B).

Note that these estimates rely on averaged quantities. In an actual device, the errors may be far stronger than estimated here due to individual components deviating from the expected performance, e.g., a shuttling lane with a local dip of the valley splitting. Such a component can be avoided by adapting the shuttling sequence in order to minimize its harmful effects at the cost of detours and additional idling [35].

B. Context and comparison with superconducting qubits

We finally assess the utility of the spin qubit device for near-term algorithms. Proposition 2 of Ref. [51] states that there is a maximum depth for a quantum optimization algorithm that consists of a fraction f_1 (f_2) of single- (two-) qubit gate layers with local error probability p_1 (p_2): Above a total of

$$D_{\max} = \frac{\log \epsilon^{-1}}{2(f_1 p_1 + f_2 p_2)} \quad (26)$$

layers, there is a classical algorithm that can sample from a Gibbs state in polynomial time, with an error $\epsilon \|H_{\text{op}}\|$ with respect to the output of the noisy quantum algorithm. Assuming that the errors are dominated by shuttling and idling, rather than gate errors—which is justified by the observations in Fig. 8—we estimate $p_{1(2)}$ by assuming that the errors are equally distributed over the shuttling and idling time. Then we weight them by the share of single- and two-qubit layers of the total distance and duration, as discussed in Sec. III A.

For the spin bus architecture with optimal shuttling velocity, we find that with $\epsilon = 0.1$ there are $D_{\max} \gtrsim 280$ gate layers possible before the quantum advantage is lost. While this estimated circuit depth is comparable with or less than what other hardware platforms with stationary qubits currently promise [46–48, 51], we emphasize that this corresponds to $p \approx 31$ rounds of parity QAOA independent of the system size. In more conventional architectures, where layers of SWAP gates are used to emulate the required connectivity between stationary qubits, the number of layers scales with the system size, thus with $O(10^2)$ qubits $D_{\max}$ allows only a few rounds of QAOA [50]. This astounding result is achieved by abstracting the problem to the parity architecture, which naturally fits the spin bus topology. In the following, we discuss further options for quantum error mitigation in order to recover the noisy output state.

For a more substantive comparison, we adapt the performance analysis to superconducting transmon qubits. For technical convenience, we assume a chip layout that matches the topology of the modular spin qubit architecture, where the interaction between the qubits is mediated by capacitive coupling, as opposed to nearest-neighbor exchange and spin shuttling. The resulting two-dimensional grid is composed of square and octagonal tiles alternately. The capacitive coupling allows for native CZ gates. The implementation of parity QAOA for transmon qubits is analogous to the circuit presented in Sec. III B, omitting the shuttling steps, which makes a comparison between the two hardware platforms straightforward. Note that we expect a square lattice of transmon qubits to perform slightly better due to the additional SWAP

gates required by the modular architecture, as discussed in Sec. III B.

The gate fidelity for transmon qubits can reach up to ≈ 0.9999 for single-qubit gates and 0.998 – 0.999 for two-qubit gates [61, 62], although the in-system performance of simultaneous two-qubit gates is typically lower and can be $\approx 0.996\%$ [63, 64]. We model this with depolarization channels for both single- and two-qubit gates with error probabilities of $p_{d,1q} = 3 \times 10^{-4}$ and $p_{d,2q} = (1 \times 10^{-3}) - (5 \times 10^{-3})$, respectively. These gates can be performed within $T_{1q} \approx 20$ ns and $T_{2q} \approx 50$ ns. Realistic values for decoherence and relaxation times in current transmon devices are $T_2 \approx 100$ μ s and $T_1 \approx 115$ μ s [62, 65, 66]. We assume a readout and initialization fidelity of $F_m = 0.995$ and we assume that all qubits can be read out simultaneously [65]. Compared to the spin qubits, the gates of the transmons are faster and have similar gate fidelity (slightly lower for in-system performance) and with a similar T_2 , and no shuttling is required.

We estimate the performance of the transmon chip in analogy to the spin qubits. With optimistic assumptions for the two-qubit gate fidelity, we find that the error can be between $\epsilon \approx 0.043$ for a choice of $p_{d,2q} = 2 \times 10^{-3}$ and $\epsilon \approx 0.027$ for a choice of $p_{d,2q} = 10^{-3}$, corresponding to $F_{\text{CZ}} \approx 0.9987$ and $F_{\text{CZ}} \approx 0.9993$, respectively. This result is slightly better than the optimal outcome for modular spin qubits, although of a comparable magnitude. We note that the optimal result obtained from the spin bus, whose topology naturally matches the parity architecture, lies well within the range defined by the optimistic transmon chip.

Using the more conservative estimate of $p_{d,2q} = 5 \times 10^{-3}$, corresponding to $F_{\text{CZ}} \approx 0.9967$, which was observed for simultaneous two-qubit gates integrated in a two-dimensional grid of 67 qubits [63, 64], we find $\epsilon \approx 0.087$. Both spin qubit architectures can reach and exceed this performance. The spin bus with optimized shuttling velocity can outperform this transmon chip for all considered distributions of the valley splitting, leaving a wide margin for the optimization of the semiconductor heterostructure.

In the limit of low-frequency noise resulting in a Gaussian decay of the coherences, Appendix B, both spin qubit architectures can outperform the transmon qubits even for the optimistic choice of gate fidelities. The superconducting qubit platform is hardly affected by the choice of exponential or Gaussian decay, since the execution time of the algorithm is much faster than T_2 , with gate errors being the limiting factor.

VI. DECODING AND ERROR MITIGATION

In the previous section, we assessed the expected single-qubit error probabilities for one round of QAOA both with the spin bus and a modular architecture. To evaluate the total algorithm performance, one needs to consider that parity QAOA is typically evaluated by measuring all qubits and then reconstructing the logical state from so-called spanning trees—subgraphs of $N - 1$ physical qubits connected by exactly one path that spans over all N logical qubits [36]. One such spanning tree is sufficient for reconstructing the logical state up to a global spin-flip. In the error-free case, all spanning trees yield the same logical state, which corresponds

to the output state of the algorithm. If, however, physical errors occur, different logical states are obtained from different spanning trees. In that case, the logical state with the lowest energy is accepted as the optimal result. Considering the error probabilities computed in the previous section, both architectures operated with optimal shuttling velocity can reach the low-error regime where parity QAOA has a clear advantage over conventional QAOA: Simulations of small noisy systems have shown that parity QAOA equipped with the classical postprocessing of the spanning trees can still have a high success probability, even exceeding the success probability of standard QAOA with an ideal system [36].

This is sufficient for the treatment of optimization problems, although it does not allow for a decision on whether the accepted result is the output of the algorithm or was produced by random noise. For some instances, it may also be beneficial to decode the readout results in a way that allows the reconstruction of the output state, which is particularly crucial in view of future applications of universal parity quantum computing [37]. We continue to evaluate the performance with respect to those two aspects. In Ref. [67], an estimate is given for an upper bound for the probability of the decoding to fail and result in the acceptance of the wrong logical state from the parity architecture.

This upper bound for the decoding error probability decays exponentially with the number of logical qubits. Based on the estimated wiring complexity, we assume that a spin bus processor with up to 50 unit cells can be operated with room-temperature controls [4], which allows for up to 20 logical qubits. Together with the physical error probability $\varepsilon \approx 0.037$ ($\varepsilon \approx 0.069$) of the spin bus (modular) architecture, this results in a decoding error probability of $\leq 2.5\%$ ($\leq 13.2\%$) with realistic gate errors. This represents an upper bound only, and the actual probability is expected to be much lower with a sophisticated decoding scheme. For example, when applying belief propagation, the decoding error probability can be expected to be below 1% for 6 (8) or more logical qubits, i.e., 15 (28) physical qubits [67].

For the spanning tree readout, we expect that, with a reliable algorithm, the correct output state is observed more frequently than random states. This can be exploited to decide whether the accepted result is also the correct output. To do so, we assume that n spanning trees of $N - 1$ qubits are read out and used for decoding. We also assume that they are distributed evenly over the chip in order to access all information, such that all qubits are included in one tree before any qubit is included in a second. With N logical and thus $K = N(N - 1)/2$ physical qubits, this means that all physical qubits are part of $\lfloor n(N - 1)/K \rfloor = \lfloor 2n/N \rfloor$ or $\lfloor 2n/N \rfloor + 1$ spanning trees. The number of qubits that are part of $\lfloor 2n/N \rfloor + 1$ trees is $n(N - 1) \bmod K$.

Thus, with one physical error on the chip, there is a chance of finding $\lfloor 2n/N \rfloor$ incorrect trees with probability $1 - \lfloor n(N - 1) \bmod K \rfloor / K$ and of finding $\lfloor 2n/N \rfloor + 1$ incorrect trees with probability $\lfloor n(N - 1) \bmod K \rfloor / K$. The expected number of incorrect spanning trees with one physical error is thus

$$\langle n_{\text{inc}} \rangle(1) = \left\lfloor \frac{2n}{N} \right\rfloor + \frac{1}{K} \lfloor n(N - 1) \bmod K \rfloor. \quad (27)$$

Each further error will also introduce $\langle n_{\text{inc}} \rangle(1)$ incorrect spanning trees, but has an increasing chance to affect trees that

already return the incorrect logical state. Thus, the expected number of incorrect results can be computed recursively for the m th error

$$\langle n_{\text{inc}} \rangle(m) = \langle n_{\text{inc}} \rangle(m - 1) + \max \left[1 - \frac{\langle n_{\text{inc}} \rangle(m - 1)}{n}, 0 \right] \langle n_{\text{inc}} \rangle(1), \quad (28)$$

where the max was included in order to avoid artifacts from the discrete computation. Consequently, with m physical errors, on average

$$\langle n_{\text{ok}} \rangle(m) = n - \langle n_{\text{inc}} \rangle(m) \quad (29)$$

spanning trees can be expected to return the correct logical state.

Inverting $\langle n_{\text{ok}} \rangle(m)$ and assuming that the physical errors are independent and equally distributed allows us to find the probability distribution

$$P_{\varepsilon, K}(\langle n_{\text{ok}} \rangle) = B[m(\langle n_{\text{ok}} \rangle) | \varepsilon, K] \quad (30)$$

of $\langle n_{\text{ok}} \rangle$, where $B(m | \varepsilon, K)$ is a binomial distribution of the number of errors m with the physical error probability ε and qubit number K . We now define $\langle n_{\text{ok}} \rangle_x$ such that in a fraction x of the experimental runs more than $\langle n_{\text{ok}} \rangle_x$ spanning trees returning the correct output state can be expected:

$$\sum_{\langle n_{\text{ok}} \rangle=0}^{\langle n_{\text{ok}} \rangle_x} P_{\varepsilon, K}(\langle n_{\text{ok}} \rangle) = 1 - x. \quad (31)$$

A logical state is accepted once it is obtained from more than $\langle n_{\text{ok}} \rangle_x$ spanning trees, where x parametrizes the expectation that a fraction x of the runs should allow for a positive decision to be made. Choosing a smaller x will result in a lower rate of erroneous positive decisions at the cost of requiring more runs.

As a measure of the reliability of the algorithm, we compute the probability for obtaining at least $\langle n_{\text{ok}} \rangle_x$ spanning trees returning an identical logical state from an entirely random outcome of the algorithm ($\varepsilon = 0.5$):

$$p_{\text{fail}} = \sum_{\langle n_{\text{ok}} \rangle=\langle n_{\text{ok}} \rangle_x}^n P_{0.5, K}(\langle n_{\text{ok}} \rangle). \quad (32)$$

This estimate can be viewed as a statistical test based on the probability distribution $P_{\varepsilon, K}$, whether the output state could plausibly arise from a random result.

The probability p_{fail} of accepting a random state as algorithm output, regardless of it being the best solution to the problem or not, is shown in Fig. 9 for the example of the spin bus architecture with different shuttling velocities, $x = 0.9$ and the choices of $n = N$ and $n = 2N$. Naturally, p_{fail} decreases with the qubit number because the probability for repeatedly observing a random result decreases fast with system size. The curves may show a small jump due to the discrete nature of the spanning trees, since a change in $\langle n_{\text{ok}} \rangle_{0.9}$ by 1 can be significant for a low number of logical qubits; this can be seen in the solid orange and dashed black curves.

Some curves with a high ε furthermore show a big jump to $p_{\text{fail}} = 1$. This can be explained by the fact that, for a too large error ε and a too strict x , $\langle n_{\text{ok}} \rangle_x$ will also decrease when the number of qubits is increased and may fall to 0 such that no decoding is possible. In that case, the decoding is not

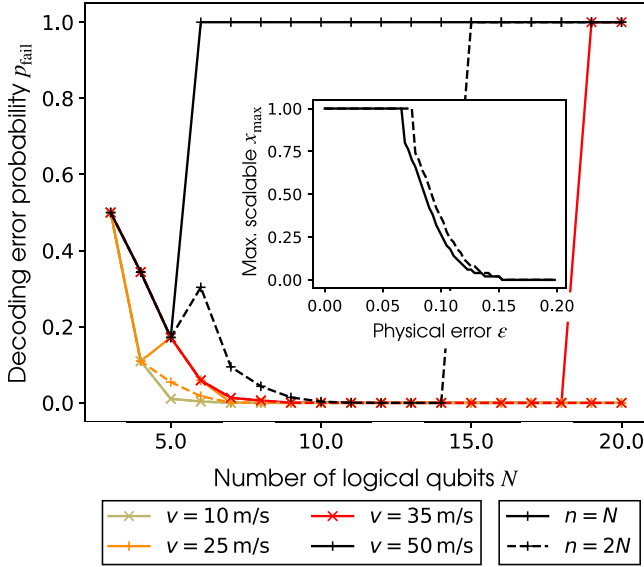


FIG. 9. Decoding based on the readout of spanning trees for up to 20 logical qubits. Main panel: Probability p_{fail} of accepting a random state as algorithm output computed with $x = 0.9$ for the spin bus as a function of the number of logical qubits, with $\bar{E}_v = 100 \mu\text{eV}$, $\sigma_{E_v}^2 = (20 \mu\text{eV})^2$. A number of $n = N$ ($2N$) spanning trees are decoded in the solid (dashed) curves. Increasing the system typically suppresses the decoding error; however, small jumps due to the discrete nature of spanning trees may occur (orange curve), and for a too large physical error the decoding may not be scalable. The requirements can be lowered by increasing the number n of spanning trees. The lines only serve as a guide to the eye. Inset: Maximal fraction x of experiments, x_{max} , that allows a decision about which state is the correct output state for any system size (scalable decoding) as a function of the error probability ϵ . If $x = 1$, the output state is reliably recovered from the decoding of the spanning trees for all qubit numbers; if $x = 0$, no decision is possible on whether the output is random or not, from a certain system size upwards. In between, decoding is still possible in general but works only for a fraction x of all attempts.

scalable. By increasing the number of spanning trees, n , more redundant information is used and thus higher error rates can be tolerated.

In the inset of Fig. 9, x_{max} is shown, which is the maximal x where $\langle n_{\text{ok}} \rangle_x$ does not decrease with the qubit number, as a function of the physical error probability. For $x_{\text{max}} = 1$ it is possible to reliably recover the output state of the algorithm, and increasing the number of qubits improves decoding. For $x_{\text{max}} = 0$, the decoding may work for small systems but will fail at a certain qubit number, that is, the algorithm will still produce a candidate solution of the optimization problem, but it is impossible to decide whether the result was produced by the algorithm or noise. In the intermediate regime, the decoding still works as usual, and increasing the system size improves the error mitigation capabilities, but also introduces a finite probability that no decision can be made on the output of the algorithm. Based on the results of Sec. V A, parity QAOA can be performed and decoded reliably with the spin bus and the modular architecture with an optimistic value for T_2 , achieved by dynamical decoupling.

VII. SUMMARY AND CONCLUSIONS

In this paper, we investigated the implementation and performance of the parity QAOA algorithm on two different electron spin qubit architectures, one based on electrons sparsely distributed over intersecting shuttling lanes (spin bus) and one where 2×2 arrays of QDs form a lattice of registers interconnected by shuttling. We presented shuttling and gate sequences for realizing all elements of parity QAOA on both platforms. While straightforward for the spin bus, the modular chip layout discussed here requires 10 successive SWAP gates to achieve the required connectivity. Alternatively, SWAP gates can be traded for hopping between adjacent QDs.

To consider realistic errors, we developed a model that allows for estimating the mean shuttling error as a function of the probability distribution of the valley splitting and the valley phase. Assuming Si/SiGe as host material and conveyor mode shuttling with realistic parameters for gate error rates, dephasing, and valley splitting, we find that both architectures can complete one round of parity QAOA with a low single-qubit error probability. The spin bus slightly outperforms the modular architecture due to the need for additional SWAP gates in the latter. The performance delivered by both platforms for parity QAOA after optimizing the shuttling process can exceed the results expected from typical superconducting transmon qubits. This result not only points out the synergy between the parity architecture and spin qubit hardware but also shows that spin qubits can be serious contenders among the leading quantum computing platforms.

Our error analysis suggests that the main limitations are dephasing from intervalley transitions if the shuttling is nonadiabatic and dephasing from environmental noise if the shuttling is too slow. Engineering the Si/SiGe interface [27–29] and Si quantum well [30,31] to deterministically enhance the valley splitting or adjusting the shuttling path and velocity to avoid excitations can further reduce the minimal error by allowing for a higher shuttling velocity. Protecting the spins from charge noise and prolonging their lifetime by dynamical decoupling can make the algorithm more robust at low velocities as well.

Finally, we discussed the possibilities for quantum error mitigation based on the results of the error model. While parity QAOA has an intrinsic error mitigation capability, we show that it is not guaranteed that the actual outcome of the algorithm can be identified. Nevertheless, our estimations indicate that the physical errors of both spin qubit platforms are low enough to decode the final state with a high success probability. This is accomplished either with dedicated decoding schemes such as belief propagation or with the classical postprocessing of spanning trees that is commonly used in conjunction with parity QAOA.

Other studies have shown that a direct implementation of QAOA requires much higher gate fidelities than currently available in any platform to achieve a quantum advantage for problems whose coupling topology does not match that of the hardware [46–48,50]. The combination of the parity encoding with an architecture that is well-matched to its requirement paints a more optimistic picture. The cost for this advantage is a quadratic overhead in the number of qubits. Importantly, this

increase in qubit number can be leveraged without requiring a higher fidelity, implying a much better scalability of parity encoded QAOA problems.

Our results highlight that a two-dimensional spin qubit platform can indeed serve as a natural implementation of the parity architecture even if the connectivity does not directly correspond to a square-lattice geometry. Thus, the utility of spin qubits for quantum computing tasks such as solving optimization problems can be advanced by alleviating the need for long-range interaction and by allowing for constant-depth QAOA with the possibility of quantum error mitigation. The results presented here show that parity QAOA may be in reach of near-term spin qubit devices, promising a substantial quantum advantage for a large class of problems once a qubit number on the order of a few thousand is achieved. It can be expected that future improvements of the qubit coherence make universal parity quantum computing viable. This will provide an advantage by reducing the circuit depth of cornerstone quantum algorithms such as the quantum Fourier transform, while relying exclusively on nearest-neighbor interactions and single-qubit gates [37].

We note that we considered only a minimal modular architecture, the performance with larger registers is an open question to be addressed in the future. Other promising directions are the use of resonator [68–72] and RF readout [73–76] of spin qubits for measurement-based parity quantum computing [41] and the use of a hybrid formulation of parity QAOA which reduces the number of constraints that are enforced explicitly and allows a modularization of the code [52]. The latter can be useful for the evolution of the modular design with larger registers. For future research in order to unlock the full potential of spin qubits equipped with the parity architecture, it is also relevant to explore whether electron or hole spin qubits can provide special advantages, and it may be beneficial to revisit and further optimize the native gate set of the platform for the interactions required here.

ACKNOWLEDGMENTS

We thank Philipp Aumann, Adu Offei-Danso, Javad Kazemi, Enrique Naranjo Bejarano, and Anette Messinger for helpful discussions, and Kilian Ender for critical comments on the manuscript. This study was supported by NextGenerationEU via FFG and Quantum Austria (FFG Project No. FO999896208). This research was funded in part by the Austrian Science Fund (FWF) under Grants DOI 10.55776/F71 and DOI 10.55776/Y1067, by the German Research Foundation (DFG) under Germany's Excellence Strategy–Cluster of Excellence Matter and Light for Quantum Computing (ML4Q) EXC 2004/1-390534769, and by the Federal Ministry of Education and Research (Germany), Funding Ref. No. 13N15652. This project was funded via Project Si-QuBus within the QuantERA ERA-NET Cofund in Quantum Technologies and within the QuantERA II Programme that has received funding from the European Union's Horizon 2020 research and innovation program under Grant Agreement No. 101017733.

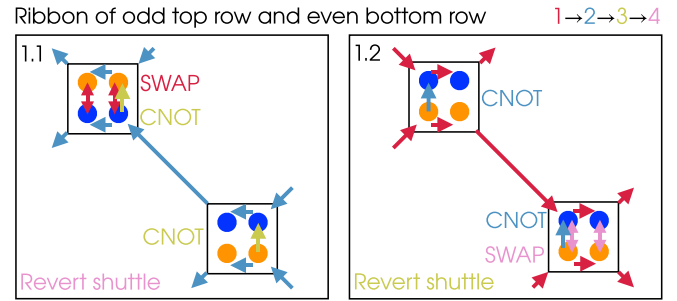


FIG. 10. Alternative circuit for implementing the parity constraints on ribbons of an odd top and an even bottom row in a modular architecture if executed on each unit cell in parallel. This sequence can replace the steps 1.1 and 1.2 of Fig. 7, and its reverse can replace step 4 of Fig. 7. Effectively, some SWAP gates are replaced by transitions between the QDs of a module, indicated by horizontal arrows. Thus, four SWAPs in steps 1 and 4 each can be avoided. The remaining SWAP gates are required for the subsequent steps.

APPENDIX A: ALTERNATIVE SEQUENCE FOR THE MODULAR ARCHITECTURE

For the modular architecture, the parity constraints can also be realized with a circuit where eight SWAP gates are traded for coherent transitions between the QDs within a unit cell and thus requires only two SWAP gates per qubit. This alternative sequence is shown in Fig. 10. The amount of shuttling operations is the same, with the exception that the electrons are now shuttled to vacant dots within the neighboring unit cells instead of adjacent sites, and the other gates remain unchanged. Thus, this approach can be a beneficial alternative if the fidelity of one intramodule transfer outperforms a SWAP gate.

With our assumptions, the fidelity of a single SWAP is expected to be $\gtrsim 99.6\%$. The transport within the module can be realized by utilizing the plunger and barrier gates for conveyor mode-shuttling or by phase-coherent bucket brigade shuttling in a double quantum dot. A fidelity per hop approaching this threshold was demonstrated in SiMOS devices [77,78]. In Si/SiGe, the probability of spin-flip errors has been shown to be in the required range for high-fidelity shuttling [79], and theoretical estimates predict a fidelity above the threshold to make this alternative viable [20].

APPENDIX B: PERFORMANCE IN THE PRESENCE OF LOW-FREQUENCY NOISE

The dominating source of dephasing in semiconductor spin qubits is considered to be charge noise, electric field fluctuations with a power spectral density $S(\omega) \propto 1/\omega^\alpha$, $\alpha \approx 1$ [1]. The low-frequency noise gives rise to a Gaussian decay of the coherences. Compared to the case of rapid fluctuations, this leads to an improved performance of the shallow algorithm discussed here for instances in which the passive dephasing is the limiting factor. We investigate the case of quasistatic noise by estimating the average single-qubit error probability ε with $p_{\phi, \text{idl}} = (t/T_2)^2$ as probability for the dephasing channel during the idling of the qubits. The results are depicted in Fig. 11.

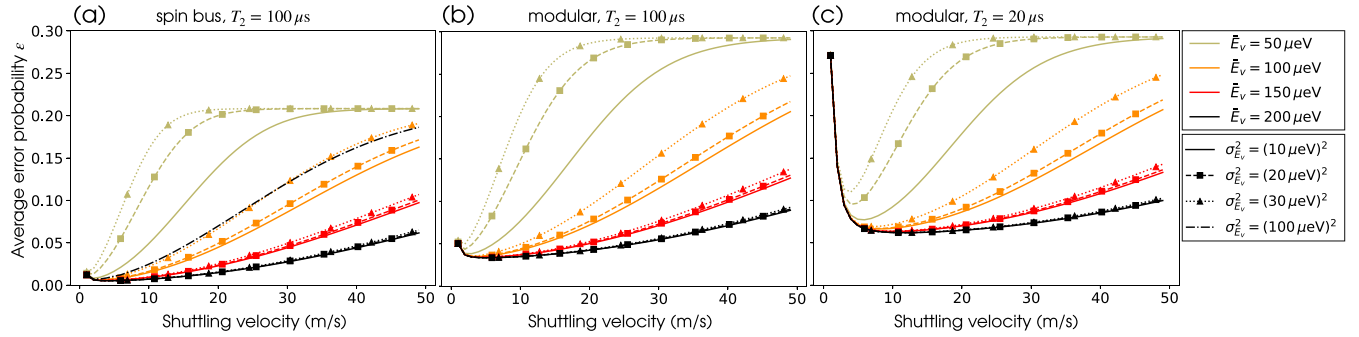


FIG. 11. Average error probability after one round of QAOA with $p_{\phi, \text{idl}} = (t/T_2)^2$ for (a) the spin bus with $T_2 = 100 \mu\text{s}$ and (b) for the modular architecture with an optimistic $T_2 = 100 \mu\text{s}$ and (c) with a realistic $T_2 = 20 \mu\text{s}$. The examples plotted here correspond to Fig. 8 in the case of quasistatic noise. The reduced dephasing during short idling times improves the performance for slow shuttling, while the dephasing due to nonadiabatic effects at high v is unaffected.

As a consequence, the error probability for a low shuttling velocity v is considerably reduced. Thus, the optimal result is found for slower shuttling and with a lower minimal error probability. Towards higher shuttling velocity, where the errors are dominated by nonadiabatic effects, the effect

vanishes. With a mean valley splitting of $\bar{E}_v \approx 200 \mu\text{eV}$ and a standard deviation of $\sigma_{E_v} \approx 30 \mu\text{eV}$ or less, a single-qubit error around $\varepsilon \approx 0.005$ ($\varepsilon \approx 0.034$) is observed for the spin bus (modular) architecture.

- [1] G. Burkard, T. D. Ladd, A. Pan, J. M. Nichol, and J. R. Petta, Semiconductor spin qubits, *Rev. Mod. Phys.* **95**, 025003 (2023).
- [2] L. M. K. Vandersypen, H. Bluhm, J. S. Clarke, A. S. Dzurak, R. Ishihara, A. Morello, D. J. Reilly, L. R. Schreiber, and M. Veldhorst, Interfacing spin qubits in quantum dots and donors—hot, dense, and coherent, *npj Quantum Inf.* **3**, 34 (2017).
- [3] O. Crawford, J. R. Cruise, N. Mertig, and M. F. Gonzalez-Zalba, Compilation and scaling strategies for a silicon quantum processor with sparse two-dimensional connectivity, *npj Quantum Inf.* **9**, 13 (2023).
- [4] M. Künne, A. Willmes, M. Oberländer, C. Gorjaew, J. D. Teske, H. Bhardwaj, M. Beer, E. Kammerloher, R. Otten, I. Seidler, R. Xue, L. R. Schreiber, and H. Bluhm, The spinbus architecture: Scaling spin qubits with electron shuttling, *Nat. Commun.* **15**, 4977 (2024).
- [5] A. M. J. Zwerver, T. Krähenmann, T. F. Watson, L. Lampert, H. C. George, R. Pillarisetty, S. A. Bojarski, P. Amin, S. V. Amitonov, J. M. Boter, R. Caudillo, D. Correas-Serrano, J. P. Dehollain, G. Droulers, E. M. Henry, R. Kotlyar, M. Lodari, F. Lüthi, D. J. Michalak, B. K. Mueller *et al.*, Qubits made by advanced semiconductor manufacturing, *Nat. Electron.* **5**, 184 (2022).
- [6] X. Xue, M. Russ, N. Samkharadze, B. Undseth, A. Sammak, G. Scappucci, and L. M. K. Vandersypen, Quantum logic with spin qubits crossing the surface code threshold, *Nature (London)* **601**, 343 (2022).
- [7] A. R. Mills, C. R. Guinn, M. J. Gullans, A. J. Sigillito, M. M. Feldman, E. Nielsen, and J. R. Petta, Two-qubit silicon quantum processor with operation fidelity exceeding 99%, *Sci. Adv.* **8**, eabn5130 (2022).
- [8] A. Noiri, K. Takeda, T. Nakajima, T. Kobayashi, A. Sammak, G. Scappucci, and S. Tarucha, Fast universal quantum gate above the fault-tolerance threshold in silicon, *Nature (London)* **601**, 338 (2022).
- [9] E. J. Connors, J. J. Nelson, H. Qiao, L. F. Edge, and J. M. Nichol, Low-frequency charge noise in Si/SiGe quantum dots, *Phys. Rev. B* **100**, 165305 (2019).
- [10] E. J. Connors, J. Nelson, L. F. Edge, and J. M. Nichol, Charge-noise spectroscopy of Si/SiGe quantum dots via dynamically-decoupled exchange oscillations, *Nat. Commun.* **13**, 940 (2022).
- [11] W. Huang, C. H. Yang, K. W. Chan, T. Tanttu, B. Hensen, R. C. C. Leon, M. A. Fogarty, J. C. C. Hwang, F. E. Hudson, K. M. Itoh, A. Morello, A. Laucht, and A. S. Dzurak, Fidelity benchmarks for two-qubit gates in silicon, *Nature (London)* **569**, 532 (2019).
- [12] I. Heinz and G. Burkard, Crosstalk analysis for single-qubit and two-qubit gates in spin qubit arrays, *Phys. Rev. B* **104**, 045420 (2021).
- [13] I. Heinz and G. Burkard, Crosstalk analysis for simultaneously driven two-qubit gates in spin qubit arrays, *Phys. Rev. B* **105**, 085414 (2022).
- [14] I. Heinz, A. R. Mills, J. R. Petta, and G. Burkard, Analysis and mitigation of residual exchange coupling in linear spin qubit arrays, *Phys. Rev. Res.* **6**, 013153 (2024).
- [15] J. M. Taylor, H. A. Engel, W. Dür, A. Yacoby, C. M. Marcus, P. Zoller, and M. D. Lukin, Fault-tolerant architecture for quantum computation using electrically controlled semiconductor spins, *Nat. Phys.* **1**, 177 (2005).
- [16] I. Seidler, T. Struck, R. Xue, N. Focke, S. Trellenkamp, H. Bluhm, and L. R. Schreiber, Conveyor-mode single-electron shuttling in Si/SiGe for a scalable quantum computing architecture, *npj Quantum Inf.* **8**, 100 (2022).
- [17] H. Flentje, P. A. Mortemousque, R. Thalineau, A. Ludwig, A. D. Wieck, C. Bäuerle, and T. Meunier, Coherent

- long-distance displacement of individual electron spins, *Nat. Commun.* **8**, 501 (2017).
- [18] T. Fujita, T. A. Baart, C. Reichl, W. Wegscheider, and L. M. K. Vandersypen, Coherent shuttle of electron-spin states, *npj Quantum Inf.* **3**, 22 (2017).
- [19] A. R. Mills, D. M. Zajac, M. J. Gullans, F. J. Schupp, T. M. Hazard, and J. R. Petta, Shuttling a single charge across a one-dimensional array of silicon quantum dots, *Nat. Commun.* **10**, 1063 (2019).
- [20] F. Ginzl, A. R. Mills, J. R. Petta, and G. Burkard, Spin shuttling in a silicon double quantum dot, *Phys. Rev. B* **102**, 195418 (2020).
- [21] V. Langrock, J. A. Krzywda, N. Focke, I. Seidler, L. R. Schreiber, and L. Cywiński, Blueprint of a scalable spin qubit shuttle device for coherent mid-range qubit transfer in disordered Si/SiGe/SiO₂, *PRX Quantum* **4**, 020305 (2023).
- [22] T. Struck, M. Volmer, L. Visser, T. Offermann, R. Xue, J.-S. Tu, S. Trellenkamp, Ł. Cywiński, H. Bluhm, and L. R. Schreiber, Spin-EPR-pair separation by conveyor-mode single electron shuttling in Si/SiGe, *Nat. Commun.* **15**, 1325 (2024).
- [23] R. Xue, M. Beer, I. Seidler, S. Humpohl, J.-S. Tu, S. Trellenkamp, T. Struck, H. Bluhm, and L. R. Schreiber, Si/SiGe QuBus for single electron information-processing devices with memory and micron-scale connectivity function, *Nat. Commun.* **15**, 2296 (2024).
- [24] B. Koiller, X. Hu, and S. Das Sarma, Exchange in silicon-based quantum computer architecture, *Phys. Rev. Lett.* **88**, 027903 (2001).
- [25] M. Friesen and S. N. Coppersmith, Theory of valley-orbit coupling in a Si/SiGe quantum dot, *Phys. Rev. B* **81**, 115324 (2010).
- [26] A. L. Saraiva, M. J. Calderón, X. Hu, S. Das Sarma, and B. Koiller, Physical mechanisms of interface-mediated intervalley coupling in Si, *Phys. Rev. B* **80**, 081305(R) (2009).
- [27] B. P. Wuetz, M. P. Losert, S. Koelling, L. E. A. Stehouwer, A.-M. J. Zwerver, S. G. J. Philips, M. T. Mądzik, X. Xue, G. Zheng, M. Lodari, S. V. Amitonov, N. Samkharadze, A. Sammak, L. M. K. Vandersypen, R. Rahman, S. N. Coppersmith, O. Moutanabbir, M. Friesen, and G. Scappucci, Atomic fluctuations lifting the energy degeneracy in Si/SiGe quantum dots, *Nat. Commun.* **13**, 7730 (2022).
- [28] J. R. F. Lima and G. Burkard, Interface and electromagnetic effects in the valley splitting of Si quantum dots, *Mater. Quantum Technol.* **3**, 025004 (2023).
- [29] J. R. F. Lima and G. Burkard, Valley splitting depending on the size and location of a silicon quantum dot, *Phys. Rev. Mater.* **8**, 036202 (2024).
- [30] T. McJunkin, E. R. MacQuarrie, L. Tom, S. F. Neyens, J. P. Dodson, B. Thorgrimsson, J. Corrigan, H. E. Ercan, D. E. Savage, M. G. Lagally, R. Joynt, S. N. Coppersmith, M. Friesen, and M. A. Eriksson, Valley splittings in Si/SiGe quantum dots with a germanium spike in the silicon well, *Phys. Rev. B* **104**, 085406 (2021).
- [31] T. McJunkin, B. Harpt, Y. Feng, P. Losert, R. Rahman, J. P. Dodson, M. A. Wolfe, D. E. Savage, M. G. Lagally, S. N. Coppersmith, M. Friesen, R. Joynt, and M. A. Eriksson, SiGe quantum wells with oscillating Ge concentrations for quantum dot qubits, *Nat. Commun.* **13**, 7777 (2022).
- [32] M. Volmer, T. Struck, A. Sala, B. Chen, M. Oberländer, T. Offermann, R. Xue, L. Visser, J.-S. Tu, S. Trellenkamp, Ł. Cywiński, H. Bluhm, and L. R. Schreiber, Mapping of valley-splitting by conveyor-mode spin-coherent electron shuttling, *npj Quantum Inf.* **10**, 61 (2024).
- [33] E. T. Campbell, B. M. Terhal, and C. Vuillot, Roads towards fault-tolerant universal quantum computation, *Nature (London)* **549**, 172 (2017).
- [34] W. Lechner, P. Hauke, and P. Zoller, A quantum annealing architecture with all-to-all connectivity from local interactions, *Sci. Adv.* **1**, e1500838 (2015).
- [35] K. Ender, R. Hoeven, B. E. Niehoff, M. Drieb-Schön, and W. Lechner, Parity quantum optimization: Compiler, *Quantum* **7**, 950 (2023).
- [36] A. Weidinger, G. B. Mbeng, and W. Lechner, Error mitigation for quantum approximate optimization, *Phys. Rev. A* **108**, 032408 (2023).
- [37] M. Fellner, A. Messinger, K. Ender, and W. Lechner, Universal parity quantum computing, *Phys. Rev. Lett.* **129**, 180503 (2022).
- [38] W. Lechner, Quantum approximate optimization with parallelizable gates, *IEEE Trans. Quantum Eng.* **1**, 1 (2020).
- [39] J. Unger, A. Messinger, B. E. Niehoff, M. Fellner, and W. Lechner, Low-depth circuit implementation of parity constraints for quantum optimization (2022), [arXiv:2211.11287](https://arxiv.org/abs/2211.11287).
- [40] M. Fellner, A. Messinger, K. Ender, and W. Lechner, Applications of universal parity quantum computation, *Phys. Rev. A* **106**, 042442 (2022).
- [41] A. Messinger, M. Fellner, and W. Lechner, Constant depth code deformations in the parity architecture, in *2023 IEEE International Conference on Quantum Computing and Engineering (QCE)* (IEEE, Piscataway, NJ, 2023).
- [42] E. Farhi, J. Goldstone, and S. Gutmann, A quantum approximate optimization algorithm (2014), [arXiv:1411.4028](https://arxiv.org/abs/1411.4028) [quant-ph].
- [43] E. Farhi and A. W. Harrow, Quantum supremacy through the quantum approximate optimization algorithm (2016), [arXiv:1602.07674](https://arxiv.org/abs/1602.07674) [quant-ph].
- [44] S. Hadfield, Z. Wang, B. O’Gorman, E. G. Rieffel, D. Venturelli, and R. Biswas, From the quantum approximate optimization algorithm to a quantum alternating operator ansatz, *Algorithms* **12**, 34 (2019).
- [45] E. Farhi, J. Goldstone, S. Gutmann, J. Lapan, A. Lundgren, and D. Preda, A quantum adiabatic evolution algorithm applied to random instances of an NP-complete problem, *Science* **292**, 472 (2001).
- [46] G. Pagano, A. Bapat, P. Becker, K. S. Collins, A. De, P. W. Hess, H. B. Kaplan, A. Kyprianidis, W. L. Tan, C. Baldwin, L. T. Brady, A. Deshpande, F. Liu, S. Jordan, A. V. Gorshkov, and C. Monroe, Quantum approximate optimization of the long-range ising model with a trapped-ion quantum simulator, *Proc. Natl. Acad. Sci. USA* **117**, 25396 (2020).
- [47] M. P. Harrigan, K. J. Sung, M. Neeley, K. J. Satzinger, F. Arute, K. Arya, J. Atalaya, J. C. Bardin, R. Barends, S. Boixo, M. Broughton, B. B. Buckley, D. A. Buell, B. Burkett, N. Bushnell, Y. Chen, Z. Chen, B. Chiaro, R. Collins, W. Courtney *et al.*, Quantum approximate optimization of non-planar graph problems on a planar superconducting processor, *Nat. Phys.* **17**, 332 (2021).

- [48] S. Ebadi, A. Keesling, M. Cain, T. T. Wang, H. Levine, D. Bluvstein, G. Semeghini, A. Omran, J.-G. Liu, R. Samajdar, X.-Z. Luo, B. Nash, X. Gao, B. Barak, E. Farhi, S. Sachdev, N. Gemelke, L. Zhou, S. Choi, H. Pichler *et al.*, Quantum optimization of maximum independent set using Rydberg atom arrays, *Science* **376**, 1209 (2022).
- [49] J. Preskill, Quantum computing in the NISQ era and beyond, *Quantum* **2**, 79 (2018).
- [50] J. Weidenfeller, L. C. Valor, J. Gacon, C. Tornow, L. Bello, S. Woerner, and D. J. Egger, Scaling of the quantum approximate optimization algorithm on superconducting qubit based hardware, *Quantum* **6**, 870 (2022).
- [51] D. S. França and R. García-Patrón, Limitations of optimization algorithms on noisy quantum devices, *Nat. Phys.* **17**, 1221 (2021).
- [52] K. Ender, A. Messinger, M. Fellner, C. Dłaska, and W. Lechner, Modular parity quantum approximate optimization, *PRX Quantum* **3**, 030304 (2022).
- [53] See Supplemental Material at <http://link.aps.org/supplemental/10.1103/PhysRevB.110.075302> for animations of the gate and shuttling sequences presented in the main text.
- [54] F. Borsoi, N. W. Hendrickx, V. John, M. Meyer, S. Motz, F. van Riggelen, A. Sammak, S. L. de Snoo, G. Scappucci, and M. Veldhorst, Shared control of a 16 semiconductor quantum dot crossbar array, *Nat. Nanotechnol.* **19**, 21 (2024).
- [55] M. A. Nielsen and I. L. Chuang, *Quantum Computation and Quantum Information* (Cambridge University Press, Cambridge, 2000).
- [56] T. Struck, A. Hollmann, F. Schauer, O. Fedorets, A. Schmidbauer, K. Sawano, H. Riemann, N. V. Abrosimov, Ł. Cywiński, D. Bougeard, and L. R. Schreiber, Low-frequency spin qubit energy splitting noise in highly purified $^{28}\text{Si}/\text{SiGe}$, *npj Quantum Inf.* **6**, 40 (2020).
- [57] T. Ando, A. B. Fowler, and F. Stern, Electronic properties of two-dimensional systems, *Rev. Mod. Phys.* **54**, 437 (1982).
- [58] P. Stano and J. Fabian, Spin-orbit effects in single-electron states in coupled quantum dots, *Phys. Rev. B* **72**, 155410 (2005).
- [59] D. V. Bulaev and D. Loss, Spin relaxation and anticrossing in quantum dots: Rashba versus Dresselhaus spin-orbit coupling, *Phys. Rev. B* **71**, 205324 (2005).
- [60] V. Srinivasa, K. C. Nowack, M. Shafiei, L. M. K. Vandersypen, and J. M. Taylor, Simultaneous spin-charge relaxation in double quantum dots, *Phys. Rev. Lett.* **110**, 196803 (2013).
- [61] Z. Li, P. Liu, P. Zhao, Z. Mi, H. Xu, X. Liang, T. Su, W. Sun, G. Xue, J.-N. Zhang, W. Liu, Y. Jin, and H. Yu, Error per single-qubit gate below 10^{-4} in a superconducting qubit, *npj Quantum Inf.* **9**, 111 (2023).
- [62] M. Kjaergaard, M. E. Schwartz, J. Braumüller, P. Krantz, J. I.-J. Wang, S. Gustavsson, and W. D. Oliver, Superconducting qubits: Current state of play, *Annu. Rev. Condens. Matter Phys.* **11**, 369 (2020).
- [63] A. Morvan *et al.*, Phase transitions in random circuit sampling (2023), [arXiv:2304.11119](https://arxiv.org/abs/2304.11119) [quant-ph].
- [64] X. Mi *et al.*, Stable quantum-correlated many-body states through engineered dissipation, *Science* **383**, 1332 (2024).
- [65] P. Jurcevic, A. Javadi-Abhari, L. S. Bishop, I. Lauer, D. F. Bogorin, M. Brink, L. Capelluto, O. Günlük, T. Itoko, N. Kanazawa, A. Kandala, G. A. Keefe, K. Krsulich, W. Landers, E. P. Lewandowski, D. T. McClure, G. Nannicini, A. Narasgond, H. M. Nayfeh, E. Pritchett *et al.*, Demonstration of quantum volume 64 on a superconducting quantum computing system, *Quantum Sci. Technol.* **6**, 025020 (2021).
- [66] A. Nersisyan, S. Poletto, N. Alidoust, R. Manenti, R. Renzas, C.-V. Bui, K. Vu, T. Whyland, Y. Mohan, E. A. Sete, S. Stanwyck, A. Bestwick, and M. Reago, *Manufacturing low dissipation superconducting quantum processors*, 2019 *IEEE International Electron Devices Meeting (IEDM)*, San Francisco, CA, USA (IEEE, 2019), pp. 31.1.1–31.1.4.
- [67] F. Pastawski and J. Preskill, Error correction for encoded quantum annealing, *Phys. Rev. A* **93**, 052325 (2016).
- [68] X. Mi, M. Benito, S. Putz, D. M. Zajac, J. M. Taylor, G. Burkard, and J. R. Petta, A coherent spin-photon interface in silicon, *Nature (London)* **555**, 599 (2018).
- [69] F. Borjans, X. Mi, and J. R. Petta, Spin digitizer for high-fidelity readout of a cavity-coupled silicon triple quantum dot, *Phys. Rev. Appl.* **15**, 044052 (2021).
- [70] R. Ruskov and C. Tahan, Quantum-limited measurement of spin qubits via curvature couplings to a cavity, *Phys. Rev. B* **99**, 245306 (2019).
- [71] B. D’Anjou and G. Burkard, Optimal dispersive readout of a spin qubit with a microwave resonator, *Phys. Rev. B* **100**, 245427 (2019).
- [72] F. Ginzler and G. Burkard, Simultaneous transient dispersive readout of multiple spin qubits, *Phys. Rev. B* **108**, 125437 (2023).
- [73] J. I. Colless, A. C. Mahoney, J. M. Hornibrook, A. C. Doherty, H. Lu, A. C. Gossard, and D. J. Reilly, Dispersive readout of a few-electron double quantum dot with fast rf gate sensors, *Phys. Rev. Lett.* **110**, 046805 (2013).
- [74] P. Pakkiam, A. V. Timofeev, M. G. House, M. R. Hogg, T. Kobayashi, M. Koch, S. Rogge, and M. Y. Simmons, Single-shot single-gate rf spin readout in silicon, *Phys. Rev. X* **8**, 041032 (2018).
- [75] A. West, B. Hensen, A. Jouan, T. Tanttu, C.-H. Yang, A. Rossi, M. F. Gonzalez-Zalba, F. Hudson, A. Morello, D. J. Reilly, and A. S. Dzurak, Gate-based single-shot readout of spins in silicon, *Nat. Nanotechnol.* **14**, 437 (2019).
- [76] M. Urdampilleta, D. J. Niegemann, E. Chanrion, B. Jadot, C. Spence, P.-A. Mortemousque, C. Bäuerle, L. Hutin, B. Bertrand, S. Barraud, R. Maurand, M. Sanquer, X. Jehl, S. De Franceschi, M. Vinet, and T. Meunier, Gate-based high fidelity spin readout in a CMOS device, *Nat. Nanotechnol.* **14**, 737 (2019).
- [77] J. Yoneda, W. Huang, M. Feng, C. H. Yang, K. W. Chan, T. Tanttu, W. Gilbert, R. C. C. Leon, F. E. Hudson, K. M. Itoh, A. Morello, S. D. Bartlett, A. Laucht, A. Saraiva, and A. S. Dzurak, Coherent spin qubit transport in silicon, *Nat. Commun.* **12**, 4114 (2021).
- [78] A. Noiri, K. Takeda, T. Nakajima, T. Kobayashi, A. Sammak, G. Scappucci, and S. Tarucha, A shuttling-based two-qubit logic gate for linking distant silicon quantum processors, *Nat. Commun.* **13**, 5740 (2022).
- [79] A. M. J. Zwerver, S. V. Amitonov, S. L. de Snoo, M. T. Mądzik, M. Rimbach-Russ, A. Sammak, G. Scappucci, and L. M. K. Vandersypen, Shuttling an electron spin through a silicon quantum dot array, *PRX Quantum* **4**, 030303 (2023).