

Hamiltonian quantum computing with superconducting qubits

A. Ciani^{1,4}, B. M. Terhal^{2,3,4}, D. P. DiVincenzo^{1,3,4}

¹ Institute for Quantum Information, RWTH Aachen University, D-52056 Aachen, Germany

² QuTech, Delft University of Technology, P.O. Box 5046, 2600 GA Delft, The Netherlands

³ Peter Grünberg Institute, Theoretical Nanoelectronics, Forschungszentrum Jülich, D-52425 Jülich, Germany

⁴ Jülich-Aachen Research Alliance (JARA), Fundamentals of Future Information Technologies, D-52425 Jülich, Germany

Abstract. We consider how the Hamiltonian Quantum Computing scheme introduced in [New Journal of Physics, vol. 18, p. 023042, 2016] can be implemented using a 2D array of superconducting transmon qubits. We show how the scheme requires the engineering of strong attractive cross-Kerr and weak flip-flop or hopping interactions and we detail how this can be achieved. Our proposal uses a new electric circuit for obtaining the attractive cross-Kerr coupling between transmons via a dipole-like element. We discuss and numerically analyze the forward motion and execution of the computation and its dependence on coupling strengths and their variability. We extend [New Journal of Physics, vol. 18, p. 023042, 2016] by explicitly showing how to construct a direct Toffoli gate, thus establishing computational universality via the Hadamard and Toffoli gate or via controlled-Hadamard, Hadamard and CNOT.

Contents

1	Introduction	2
1.1	Overview of Transmon Qubit Implementation Proposal	4
1.2	Overview of Paper	9
2	Review of the Lloyd-Terhal scheme	9
2.1	Multi-Qubit Logic	13
2.1.1	CNOT Gate	14
2.1.2	Direct Toffoli Gate	14
2.2	Dual Rail Representation	15
3	Numerical Studies of Small Lattices	16
3.1	Static Disorder	19
4	Cross-Kerr and flip-flop couplings between superconducting transmon qubits	22
4.1	Cross-Kerr interaction	22
4.2	Cross-Kerr Coupling Strength and Cross-Talk	29
4.3	Weak On-Track Flip-Flop Interactions	30
4.4	Real Interactions and Time-Reversal	32
5	Challenges	33
5.1	Loss of Excitations	33
5.2	Arrival Time and Measurement	33
6	Discussion	35
A	Hamiltonian for the direct Toffoli Gate	36
B	Feynman model with hopping particles on a 2D lattice	39
B.1	Peres' trick	42
B.2	Multi-Qubit Logic	43
C	Universality of Hadamard, CNOT and controlled-Hadamard	45

1. Introduction

In a seminal paper of the early days of quantum information, Feynman introduced a model capable of performing an arbitrary quantum computation with a time-independent Hamiltonian [1]. In this time-independent approach to quantum computing, the system is prepared in an initial state, it evolves under the action of a Hamiltonian for a certain time and is finally measured to extract the result of the computation. In Feynman's model a fundamental role is played by a quantum

clock system whose state changes upon the application of gates. While Feynman's approach to quantum computing has not received much attention at the experimental level compared to the circuit model, its importance from a theoretical point of view has long been recognized. In particular, Feynman's model has been used to analyze the QMA-completeness of the k -local Hamiltonian problem [2], and to show a formal equivalence between adiabatic and circuit-based quantum computation [3, 4].

One of the challenges in a practical implementation of the Feynman Hamiltonian is the presence of the clock system, which requires high-weight and non-local interactions in space. An alternative to the concept of a global clock is a model where each information-carrying particle has its own local clock. This idea of an asynchronous cellular automaton was first formulated by Margolus [5], and analyzed in much greater detail in Refs. [6, 7, 8, 9, 10, 11]. The idea has been formalized under the name 'space-time circuit-to-Hamiltonian' construction in Ref. [9]. In this model the position of each particle with an internal, information-carrying, state represents its local clock. In order to implement a computation involving multi-qubit gates it is necessary to achieve coordination between the local clocks, i.e. clock times need to align for particles to interact. The need for alignment requires an attractive interaction between the particles.

There are alternative constructions showing universal Hamiltonian quantum computation in which mobile multiple interacting particles do not require particles to move together. In these constructions particles interact as wavepackets in scattering regions [12, 13]. Ref. [14] has considered how to implement such alternative multi-particle walk with an ultra-cold bosonic atom system.

The appeal of a Hamiltonian approach to quantum computing is that it does not require active driving fields to enact logic: information carriers are moving through gates in space instead of time-dependent gates being applied to stationary qubits. It means that a realization of this approach is much closer to the idea of analog quantum simulation [15], but with the benefit of allowing for universal computation. In a 2D lattice realization, information is entered on the side of the lattice while the passive interactions in the bulk of the lattice are engineered to implement a chosen 1D quantum circuit, i.e. a quantum circuit with nearest-neighbor gates between qubits on a 1D line.

This computational model also lends itself well to quantum learning since gates executed in regions on the lattice can simply depend on angle parameters which can be altered from one run of the computation to the next.

In this paper we consider the scheme for Hamiltonian quantum computing on a two-dimensional lattice developed by Lloyd and Terhal in Ref. [11]. The scheme is most easily formulated in terms of particles carrying an internal spin degree of freedom, that can sit on the sites of the lattice. The lattice grid is depicted in Fig. 1. To each particle we associate a horizontal track composed of different sites on which the particle can hop. When hopping takes place a corresponding unitary gate is applied to the internal spin degree of freedom. The coordination of the motion is achieved by means of strong attractive interactions between particles on adjacent tracks. These attractive terms are associated with the edges of the lattice depicted in Fig. 1. As shown in [11] the problem

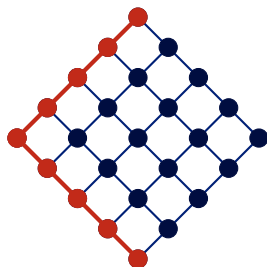


Figure 1: Example of a small rotated grid where particles can hop over $N_{\text{track}} = 9$ horizontal tracks. Red dots denote the sites that are occupied by the particles: shown is an initial configuration.

can be mapped to qubits using a dual rail encoding. In this mapping, the necessary interactions to be engineered between the qubits are strong interactions at the edges which induce excitations to move together and weaker hopping or flip-flop interactions across the plaquettes.

In this paper, we show how the scheme of Lloyd and Terhal can form the basis of a very concrete architecture using a planar array of transmon qubits. We show how the scheme requires strong attractive cross-Kerr and weak hopping interactions and how one could go about putting these together. Before going into details, we summarize the overall structure and features of our proposed architecture in the next Section 1.1.

1.1. Overview of Transmon Qubit Implementation Proposal

Consider the following layout sketch shown in Fig. 2. On each site of the lattice we have a pair of transmon qubits [16]: the ground-state $|00\rangle$ denotes the absence of a qubit information carrier, while the single-excitation states $|01\rangle$ and $|10\rangle$ represent the qubit. This means that at a certain site only one of the two transmons is in the excited state. Furthermore, on each horizontal track there is at most a single transmon excitation carrying the qubit information, all other transmon qubits are in the $|0\rangle$ state. From here onwards the word track refers to horizontal tracks.

We consider a system of grounded transmon qubits, in which each interaction is implemented by a coupling element, which might be direct or indirect, in a modular way. In addition, as a design principle, we try to design a system that is as passive as possible, meaning that we do not use active pulses or parametric drives to engineer the interactions. However, we will allow the presence of constant flux biases.

The scheme does not require any active control apart from initialization and readout and works in the following way. The feedline on the left of Fig. 2 put transmon excitations on the left sites representing some computational state $|x\rangle$ with x representing some bitstring. The system evolves under the action of an engineered Hamiltonian that implements a quantum circuit. After a certain time the qubits on the right side of the lattice are measured via a standard dispersive readout. This read-out is indicated by the read-out resonator + feedlines on the right in Fig. 2. We discuss the

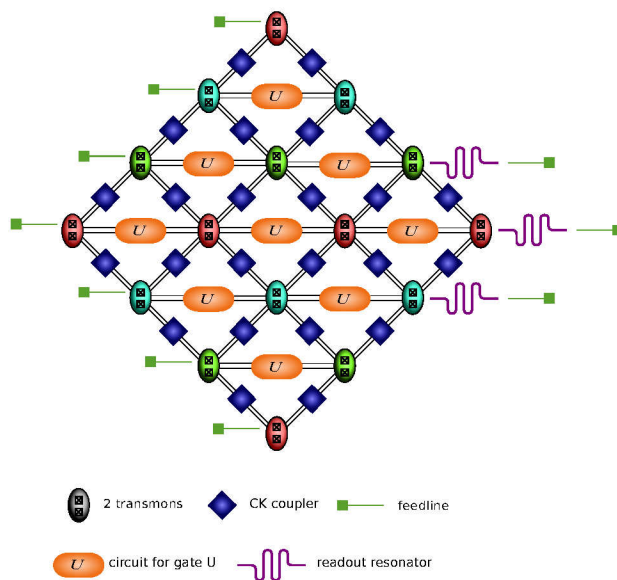
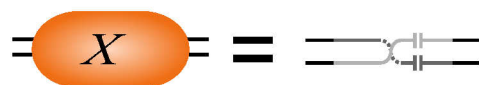
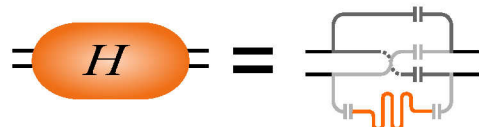


Figure 2: Global layout concept for the Hamiltonian quantum computing scheme with superconducting transmon qubits. A layout with $N_{\text{track}} = 7$ horizontal tracks is shown. The microwave feedlines on the left prepare excitations in a subset of transmons on the left (one excitation per transmon pair) to set an initial bitstring of the computation. The bulk of the grid is used to realize gates such as $U = I$, $U = H$ (Fig. 3) or $U = \text{CNOT}$, controlled-Hadamard or a Toffoli gate. If a CNOT or controlled-Hadamard is to be executed, the couplers in a region of the lattice are modified as shown in Fig. 5. The blue interactions represent doubled cross-Kerr (CK) interactions, see Fig. 4.



(a) X gate circuit.



(b) Hadamard gate circuit.

Figure 3: Examples of quantum electric circuits for implementing single-qubit gates, constructed via direct or resonator-mediated capacitive couplers.

need for measuring qubits in a larger measurement *region* and the idea of employing a MERA-like trap region in Section 5.2.

The two transmon qubits which make up a pair should be identical in their dressed frequency $\ddagger$ so that the qubit carrier space is degenerate. Furthermore, all transmon

$\ddagger$ Dressed frequency means their frequency in the presence of all couplers.

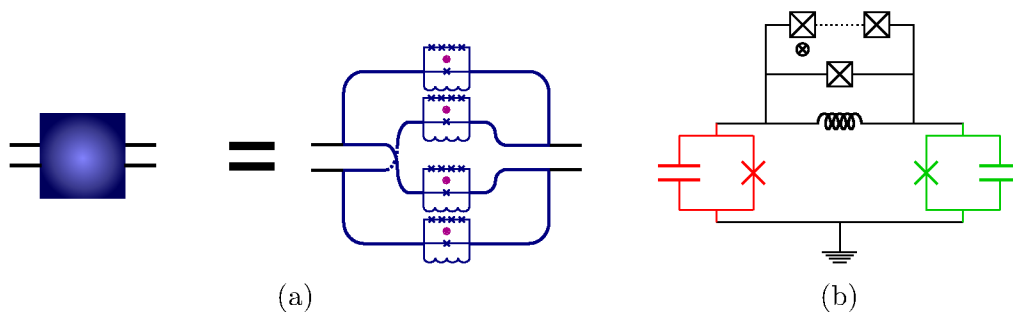


Figure 4: (a) Schematic circuit for realizing the doubled *attractive* cross-Kerr interactions between two sets of transmon pairs, one pair at each site. The two input (output) lines denote one pair. The cross-Kerr interaction is doubled in the sense that it is a sum of four cross-Kerr interactions between the pairs. (b) Coupling element that realizes the cross-Kerr interaction between two transmons on different tracks. The coupler is represented in more detail in Fig. 19, and analyzed in Sec. 4.

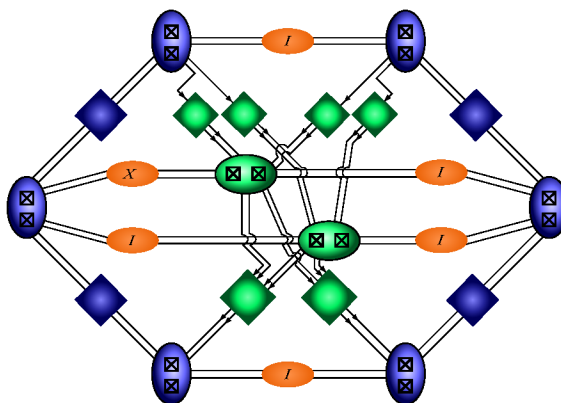


Figure 5: Coupling scheme for a CNOT region: all green edges in Fig. 9 are represented by green cross-Kerr couplers. The figure represents the two plaquettes that need to be modified in order to implement a CNOT, as detailed in Subsec. 2.1. We have a total of 16 transmons and a total of 6 elements applying unitaries like in Fig. 3. The green cross-Kerr couplers have a variable number of input and output lines: between *each* input and output line there should exist an attractive cross-Kerr interaction formed by the electric circuit in Fig. 4b. This dense coupling, in particular between the transmons on the bottom track and the middle track, makes the CNOT region the most complicated element of our proposal.

qubits on the same track have, in principle, the same dressed frequency.

Transmon qubits on adjacent tracks are detuned: in Fig. 2 we have chosen to use three frequencies, shown as three different colors, a pattern which can be repeated throughout the lattice. This detuning between qubits on adjacent tracks helps in mitigating unwanted cross-talk between the tracks. Such cross-talk is an unwanted side-effect of the circuit which induces the cross-Kerr couplings which we design to minimize. In Table 1 we list all relevant energy scales for the couplers and interactions

that are used in Fig. 2.

Flip-flop coupling J $J(\sigma_+^1 \sigma_-^2 + \sigma_-^1 \sigma_+^2)$	$2\pi \times 10 \text{ MHz}$
Transmon cross-Kerr coupling Δ $(-\Delta a^\dagger a b^\dagger b)$	$2\pi \times 100 \text{ MHz}$
Transmon freq. detuning between adj. tracks	$\gtrsim 500 \text{ MHz}$
Total transmon freq. bandwidth	$\approx 1.5 \text{ GHz}$

Table 1: Estimates of typical parameters achievable in our implementation.

All gates are generated by direct capacitive couplings between the transmon qubits and capacitive couplings via an intermediate resonator. The coupling mechanism behind all these gates are weak $J/2\pi = O(1) - O(10)$ MHz flip-flop (also referred to as hopping) couplings which move the transmon excitations along the track. In Fig. 3 we see an example of a bit-flip X and a Hadamard H gate. The Hadamard gate requires a mediated flip-flop interaction via an off-resonantly coupled resonator, so as to obtain a sign change in the coupling parameter associated with the Hadamard gate, i.e. the transition from $|1\rangle$ to $|1\rangle$ has amplitude $-1/\sqrt{2}$. By changing the strength of such flip-flop couplers, for example by tuning the frequency of the resonator, one can get any real single-qubit rotation $U(\theta)$.

The most challenging interactions to achieve with transmon qubits are the strong attractive cross-Kerr interactions on the edges. When two transmon qubits are viewed as anharmonic oscillators with annihilation operators a and b , these cross-Kerr interactions are of the form $-\Delta a^\dagger a b^\dagger b$.

One of the main problems in obtaining these interactions is that they have a tendency to come together with unwanted flip-flop and/or cross-talk interactions. These flip-flop interactions would move information-carrying excitations from one track to another which is unwanted. Secondly, two successive intertrack flip-flops can induce a logical X error by effectively moving the excitation from one transmon in a pair to the other transmon of the pair. Since we keep transmon qubits on adjacent tracks detuned, the first-order effect of residual flip-flop interactions is small. However, a sequence of two intertrack hops is a resonant process, so the strength of such terms needs to be small.

We propose a new element capable of generating large cross-Kerr interactions, while keeping the cross-talk moderate. The coupler is composed of an array of a few junctions in parallel with a smaller Josephson junction and it is schematically depicted in Fig. 4. The coupler requires a constant flux bias. Similar elements, but for (different) simulation purposes, have been proposed in Refs. [17, 18] and realized in [19], [20]. In these last papers, an attractive cross-Kerr of strength $\Delta = 2\pi \times O(7-10)$ MHz was measured which is below what we believe could be achieved with our coupler. For all such couplers, one

expects that the maximal strength of the cross-Kerr interaction will be limited by the anharmonicity of the coupled modes, see also e.g. [21], implying that anharmonic modes are needed to build cross-Kerr interactions.

In Fig. 4, we see that at each normal edge we need four couplers, which is a consequence of the fact that in the dual rail encoding we need two transmons per site. In Fig. 5 we show the couplers for a CNOT gate. The implementation of controlled unitaries, such as the CNOT, requires the splitting of an intermediate site into two sites, where the particle is routed depending on the state of the control. Thus, on the intermediate site we will have a total of four transmon qubits as can be seen from Fig. 5 (green sites).

The correct forward motion of the computation is realized in the limit $J/\Delta \ll 1$: in this limit, excitations only move when they do not incur energy penalties due to the strong cross-Kerr interaction, implying that they occupy nearest-neighbor lattice sites. The idea is that one can also selectively use the cross-Kerr interaction to determine where one excitation moves depending on the presence of another excitation, and thus what controlled state-change the qubit-carrying-excitation undergoes. This is the idea behind the CNOT, Controlled-Hadamard and Toffoli gates. The need to work in the limit $J/\Delta \ll 1$ while the computation time scales as $\sim 1/J$ requires the strongest possible Δ , in turn allowing for the largest possible J to execute the computation before excitations get lost via T_1 -relaxation. Let's assume a typical transmon $T_1 = 50 \mu s$. The computation time T_{comp} scales as $1/J \times \text{HoppingDepth}$ where HoppingDepth measures the maximum # site-to-site hoppings to execute some 1D nearest-neighbor circuit. If we consider an $N \times N$ lattice, then the probability to lose any of the transmon excitations scales as the number of excitations, namely $2N - 1$, times their probability to decay in time T_{comp} , proportional to $\approx T_{\text{comp}}/T_1$. If we fix the total loss probability to be 0.1 and take a HoppingDepth of 20 (a CNOT requires 3), one obtains $N = 15$. This rough estimate shows how much computation can be executed in our proposed scheme given current transmon relaxation times.

The ideal model does not require the doubled cross-Kerr couplings to be identical throughout the lattice. While in principle the strength Δ should be the same for all interactions between two adjacent tracks, it can be different for different pairs of tracks without affecting the motion of the computational wavefront, as long as J/Δ is small. In addition, small variability in Δ for cross-Kerr interactions between tracks may not directly harm the computation. We discuss the interesting effect of static disorder on qubit frequency, hopping and cross-Kerr interaction in Section 3 and 5.

A simpler, early-implementation, version of our model would be one without any computational logic. In this model one studies the dynamics of the transmon excitations. Each track consists of a single transmon qubit (not a pair representing a dual-rail encoded qubit) and transmon excitations would be supplied at the side of the lattice. The goal would be to experimentally observe how the combination of strong cross-Kerr interaction combined with weak hopping gives rise to a wavefront of excitations, forming a connected string which is propagating over the lattice. The engineering of this simple

Hamiltonian is close to that of a Bose-Hubbard model with strong cross-Kerr interactions on edges, self-Kerr anharmonicity for transmon qubits and weak capacitively-induced hopping across plaquettes.

1.2. Overview of Paper

The paper is organized as follows. In Sec. 2 we review the model presented in Ref. [11]. After briefly recalling the idea of the CNOT implementation in Sec. 2.1.1, we show how to construct a direct Toffoli gate in Sec. 2.1.2 (with technical details in Appendix A). In Sec. 2.2 we review the effect of the dual-rail encoding, leading to the coupling scheme shown in Fig. 2. In Sec. 3 we numerically analyze errors in the correct forward propagation of the computation including the effects of variability in couplers and frequencies for small system sizes.

In Sec. 4 we give details of the proposed implementation with transmon qubits describing the basic coupling elements that we sketched in Fig. 3 and Fig. 4. In Section 4.2 we discuss the strength of unwanted interactions mediated by the cross-Kerr couplers between three transmon qubits.

In Section 5 we discuss a challenge inherent in the realization of this computing architecture concerning the read-out of the computation. In Appendix B we discuss a 2D lattice version of the Feynman-Kitaev Hamiltonian as an alternative to the Lloyd-Terhal scheme. We show how multi-qubit gates can be done in this scheme with ideas similar to those in Sec. 2.1. Appendix C shows how to use the simpler controlled-Hadamard gate to make a Toffoli gate. We end the paper with some open questions in Sec. 6.

In all Sections but Section 4 we set $\hbar = 1$.

2. Review of the Lloyd-Terhal scheme

In this section we review some elements of the scheme for quantum computing with a time-independent Hamiltonian proposed in Ref [11]. The scheme is closely related to the one proposed in Ref [6] and the one analyzed in [10, 9]. The main difference between the scheme in Ref. [11] and the ones of Refs. [6] and [10] is the way gates are implemented. The model is most simply explained in terms of particles with spin-1/2 moving on tracks: in Section 2.2 we discuss the dual-rail representation in terms of transmon qubits. We start by considering the rotated $N \times N$ lattice in Fig. 6 with $N_{\text{track}} = 2N - 1$ tracks. A site is denoted by (i, j) , with $i, j \in \{1, 2, \dots, N\}$. Importantly, we will consider that at each time there is *only* one particle per track §. This is ensured by initializing the system in a state with this property and by ensuring that the Hamiltonian preserves it.

The quantum information is encoded in the spin degree of freedom of the particle. Particles are allowed to hop horizontally from one site to the next one (and vice-versa), –we say that the particles move on tracks–, and when hopping takes place a single-qubit

§ Note that a track is not identified by a constant index i or j , but by a constant difference $i - j$ in our notation.

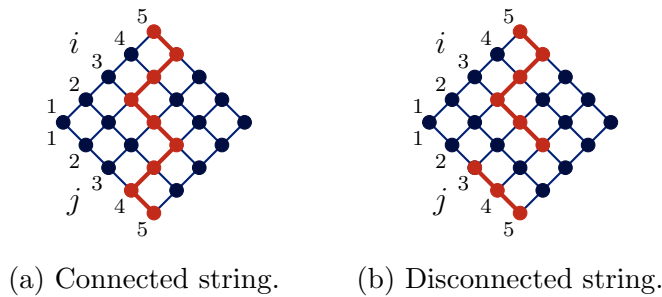


Figure 6: Examples of connected and disconnected strings of particles. The red dots denote the position of the particles.

gate can be applied to the spin degree of freedom. We can thus associate a single-qubit gate to each plaquette of the lattice in Fig. 6. If our purpose were to simulate a quantum circuit with only single-qubit gates, then independent hopping dynamics for each qubit degree of computation would clearly suffice. One way of realizing two-qubit gates is to ensure that particles move coherently together, as a string of particles [6, 10]. If we let the spin-degree of freedom of one particle influence the whereabouts and the single-qubit logic on another particle, spin-controlled single-qubit gates can be realized. An efficient forward computation can be achieved also in the context of quantum walk on a line, an idea that we will use in the discussion of the alternative lattice model in Appendix B.

The dynamics of the system is designed in such a way that if the system starts in a state in which the particles are connected as a string in Fig. 6a, it will evolve only into states in which the particles remain connected. This last property will not be guaranteed exactly by the Hamiltonian, but perturbatively. We will refer to states with connected particles as valid or connected strings. Note that if these properties are satisfied, assuming that the system starts in the connected string in which all particles are on the left of the lattice, a particular connected string univocally identifies the gates that have been executed. The particular wedge-like geometry of the rotated lattice may seem artificial when the goal is to execute a nearest-neighbour circuit. However, this geometry can be shown to guarantee an efficient forward motion of the computation on the left-half of the lattice where the number of connected strings is an increasing function of lattice depth [10, 11]. The motion of the string is also depicted in Fig. 13.

In the following part of this section we focus on the implementation of quantum circuits with only single-qubit gates, leaving the discussion of two- and three-qubit gates to the section 2.1.

Before introducing Hamiltonians, let us give some mathematical definitions. We define the particle number operator at a site (i, j) as $\mathbf{n}[i, j] = \sum_{s=0,1} n_s[i, j]$, where $n_s[i, j]$ is the number operator for particle at site (i, j) in internal spin state $s = 0, 1$. The operators $n_s[i, j]$ can be written in terms of creation and annihilation operators for a particle at site (i, j) and internal state s , i.e., $n_s[i, j] = a_s^\dagger[i, j]a_s[i, j]$. To describe the CNOT and Toffoli gate, we use split-sites labeled as (i, j, k) , $k = 0, 1$ so that $n_s[i, j, k] = a_s^\dagger[i, j, k]a_s[i, j, k]$, $\mathbf{n}[i, j, k] = \sum_{s=0}^1 n_s[i, j, k]$ and $\mathbf{n}[i, j] = \sum_{k=0}^1 \mathbf{n}[i, j, k]$.

We call the set of edges E and the set of plaquettes P . A site (i, j) will occasionally be denoted with a compact symbol μ or ν . We will identify an edge e by the sites it connects and write $e = (\mu, \nu)$.

The previously sketched ideas translate to a Hamiltonian

$$H = H_{\text{valid}} + V_{\text{hop}}. \quad (1)$$

where H_{valid} is defined as

$$H_{\text{valid}} = -\Delta \sum_{(\mu, \nu) \in E} \mathbf{n}[\mu] \mathbf{n}[\nu], \quad (2)$$

where $\Delta > 0$. Consider the spectrum of H_{valid} in the particle number basis. The groundspace of H_{valid} , in the sector with one excitation per track, is degenerate with eigenvalue $E_0 = -(N_{\text{track}} - 1)\Delta$ and composed of all possible connected strings. The number of connected strings on the lattice is $\binom{2(N-1)}{N-1}$. Thus, including the spin-degree of freedom the groundspace is $2^{N_{\text{track}}} \binom{2(N-1)}{N-1}$ -dimensional. The first excited subspace is formed by the subspace of strings which are disconnected at a single site. These strings have an energy gap of Δ above the groundspace, i.e., $E_1 - E_0 = \Delta$. In general, H_{valid} has $N_{\text{track}} - 2$ eigen-subspaces with energy $E_k - E_0 = k\Delta$, where $k \in \{0, 1, \dots, N_{\text{track}} - 1\}$ denotes the number of points where the string is broken. Note that this Hamiltonian is fully degenerate with respect to the spin degree of freedom. However, no harm is done by having additional single-site terms of the form $\omega \mathbf{n}[\mu]$ in the Hamiltonian as long as ω is the same along the entire track (as there is a only single particle per track). In addition, the attractive intertrack strengths Δ in H_{valid} can be different for one pair of tracks as compared to another pair.

The hopping Hamiltonian V_{hop} is responsible for the dynamics of the system and for the application of the gates. Let us suppose that we want to run a quantum circuit composed only of single-qubit gates. In particular, we associate a single-qubit gate, denoted by U_p , with each plaquette of the lattice, i.e., $p \in P$, for a total of $(N - 1)^2$ single-qubit gates. The hopping Hamiltonian is defined as

$$V_{\text{hop}} = -J \sum_{p \in P} V_{\text{hop},p}, \quad (3)$$

where $V_{\text{hop},p}$ is the hopping associated with the plaquette p and defined as

$$V_{\text{hop},p} = \sum_{s=0}^1 \sum_{s'=0}^1 \langle s' | U_p | s \rangle a_{s'}^\dagger [i+1, j+1] a_s [i, j] + \text{h.c.} \quad (4)$$

The overall sign of the hopping Hamiltonian V_{hop} is taken to be negative, with $J > 0$, but identical dynamics is obtained when all $J \rightarrow -J$. The effect of this hopping Hamiltonian can be exemplified by considering a single plaquette in Fig. 7a which executes the X gate when hopping takes place

$$H_X = -JV_{\text{hop},X} = -J\{a_0^\dagger[R]a_1[L] + a_1^\dagger[R]a_0[L] + \text{h.c.}\},$$

where $a_s[S]$, $s = 0, 1$, $S = L, R$ annihilates the state $|s\rangle_S$, i.e., $a_s[S] |s\rangle_S = |\text{vac}\rangle_S$.

Let us suppose that the system starts at $t = 0$ in the state $|\Psi_{\text{in}}\rangle = |\varphi_{\text{in}}\rangle_L \otimes |\text{vac}\rangle_R = (\alpha |0\rangle_L + \beta |1\rangle_L) \otimes |\text{vac}\rangle_R$. One can observe that $H_X |\Psi_{\text{in}}\rangle = -J |\Psi_{\text{out}}\rangle$ with $|\Psi_{\text{out}}\rangle = |\text{vac}\rangle_L \otimes |\varphi_{\text{out}}\rangle_R = |\text{vac}\rangle_L \otimes X_R |\varphi_{\text{in}}\rangle_R$ and $H_X |\Psi_{\text{out}}\rangle = -J |\Psi_{\text{in}}\rangle$, thus undoing the X gate.

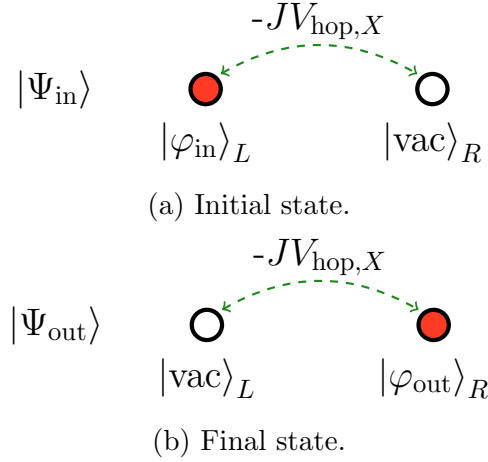


Figure 7: Example of hopping executing a single-qubit X -gate.

As long as $X = U_p$ is a unitary matrix, we see that the dynamics is not influenced by what unitary is implemented. and it is equivalent to a continuous-time quantum walk on a line with 2 sites, where $|\Psi_{\text{in}}\rangle$ and $|\Psi_{\text{out}}\rangle$ play the role of the discretized allowed positions. If we would consider a chain of $L - 1$ gates applied like in Fig. 7, in series, it can be easily shown that the dynamics is equivalent to that of a continuous-time quantum walk on a line with L sites (see also Appendix B).

Let us now examine the combined effect of V_{hop} and H_{valid} . The hopping Hamiltonian V_{hop} is not diagonal in the eigenbasis of H_{valid} and can cause transitions from valid strings to valid strings, which, in quantum optics language, are resonant. In addition, it can induce transitions from a valid connected string to an invalid, disconnected string which is a far off-resonant (by Δ), suppressed transition.

In [11] it was thus argued that a perturbative treatment of this Hamiltonian H gives rise to an effective Hamiltonian H_{eff} which only contains resonant controlled-hopping terms, i.e. particles can only move forward when their move keeps the particles in the valid string subspace. Explicitly,

$$H_{\text{eff}} = -J \sum_{p \in P} H_{\text{cond.hop},p} + \mathcal{O}(\|V_{\text{hop}}\|^2/\Delta), \quad (5)$$

where we defined the conditional hopping Hamiltonian in a plaquette p bordered by top and bottom sites resp. $(i + 1, j)$ and $(i, j + 1)$ as

$$H_{\text{cond.hop},p} = \mathbf{n}[i + 1, j] \mathbf{n}[i, j + 1] \otimes V_{\text{hop},p}. \quad (6)$$

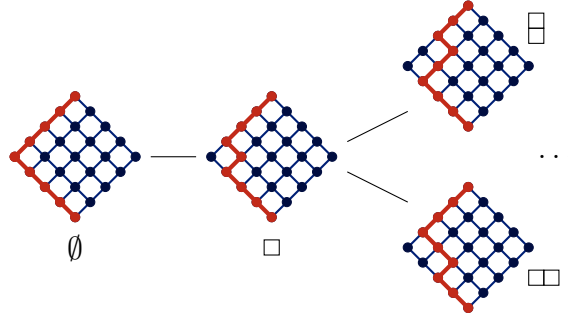


Figure 8: Coherent quantum walk of the connected string which can be viewed as a quantum walk on Young's lattice [10, 22].

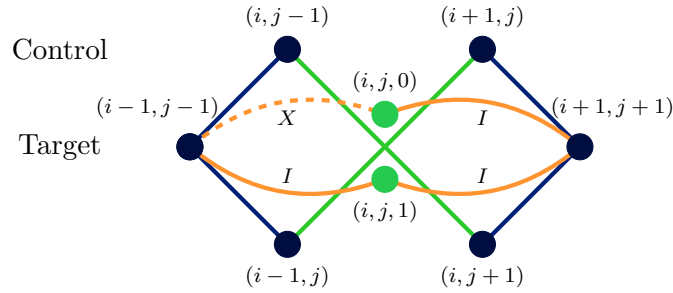


Figure 9: The CNOT gate. The control particle moves on the upper track and the target particle moves on the middle track, while a bystander particle not involved in the logic can move on the bottom track. In the middle track one has two split-sites which both connect via green edges to the other four sites. These four edges are modified from standard edges to implement the conditional logic, see Eq. 7. The particle on the middle track can go through either of the split-sites, but only if it goes past one site it undergoes an X gate.

Assuming that the system starts in the configuration in which all particles are on the left (see Fig. 8), the dynamics of the string under H_{eff} in Eq. (5) can be nicely exactly solved, see [10, 6]. In particular, the forward motion of the string in the open-wedge region, in the limit of a large lattice size, has a constant velocity [10].

2.1. Multi-Qubit Logic

To get computational universality, one could only focus on circuits involving single- and two-qubit gates, for instance Hadamard, single-qubit $T = \text{diag}(1, e^{i\pi/4})$ gate and the CNOT gate [23]. The T gate requires complex hopping parameters [11], but in our transmon implementation only real flip-flop interactions are available, see the arguments in Section 4.4. However, we can explicitly construct a way of performing the Toffoli gate achieving universality with Hadamard and Toffoli [24]. In Appendix C, we also show how universality with real gates can be achieved using Hadamard, CNOT and controlled-Hadamard gates. The controlled-Hadamard can be implemented with similar resources

as the CNOT gate.

In the following two sections we review the CNOT construction presented in [11] and introduce the construction of the Toffoli gate. We point out that these constructions may have applications that go beyond the present model, as units that implement particular gates in a modular architecture approach to quantum computation as proposed in related work [25].

2.1.1. CNOT Gate The CNOT construction is depicted in Fig. 9. The Hamiltonian of the region of the lattice where the CNOT is implemented has to be modified. In particular, the central site is replaced by two split-sites and the target particle is directed to one of the two split-sites depending on the state of the control particle. Accordingly, an X gate is applied if the control is in the $|1\rangle$ state, while the identity gate is applied if it is in the $|0\rangle$ state. Finally, the CNOT is completed via hopping of the particle from the intermediate split-sites to a final site with the application of an identity gate. This way of implementing the CNOT resembles that of a railroad switch [26]. The working principle of the CNOT also resembles that of a quantum spin transistor that has been more recently proposed in Refs. [27, 28].

How can we make sure that the logic is implemented correctly? The idea is to give energy penalties to configurations that perform incorrect logic. This means that if the control is in internal state $|0\rangle$ and the target particle is at the intermediate site $(i, j, 1)$, this invalid configuration should have a penalty Δ compared to the valid configuration in which the target particle is at $(i, j, 0)$. Analogously, a penalty is given to configurations in which the control particle has internal state $|1\rangle$ and the target particle is at site $(i, j, 0)$. This is done by letting the attractive interactions between the control particle at sites $(i, j - 1)$ and $(i + 1, j)$ and the intermediate sites $(i, j, 0)$ and $(i, j, 1)$ depend on the spin state of the control particle. The attractive interactions between the sites $(i - 1, j)$ and $(i, j + 1)$ and the split-sites are the same as before. More precisely, the green edges in Fig. 9 are chosen as

$$\begin{aligned}
 & -\Delta \sum_{s=0,1} (n_s[i, j - 1] \mathbf{n}[i, j, s] + n_s[i + 1, j] \mathbf{n}[i, j, s]) + \\
 & -\Delta \sum_{s=0,1} (\mathbf{n}[i - 1, j] \mathbf{n}[i, j, s] + \mathbf{n}[i, j + 1] \mathbf{n}[i, j, s])). \quad (7)
 \end{aligned}$$

In addition, the orange hopping edges V_{hop} in Fig. 9 to the split-sites are of the form in Eq. (4) where the X and I labels indicate whether the hopping affects the internal state.

Using the CNOT idea one can implement any controlled unitary, and in particular the controlled Hadamard, which can be used together with the CNOT to construct a Toffoli gate as shown in Appendix C.

2.1.2. Direct Toffoli Gate We now show that a Toffoli gate can be constructed in close analogy to the CNOT discussed in the previous subsection. The main idea is conveyed

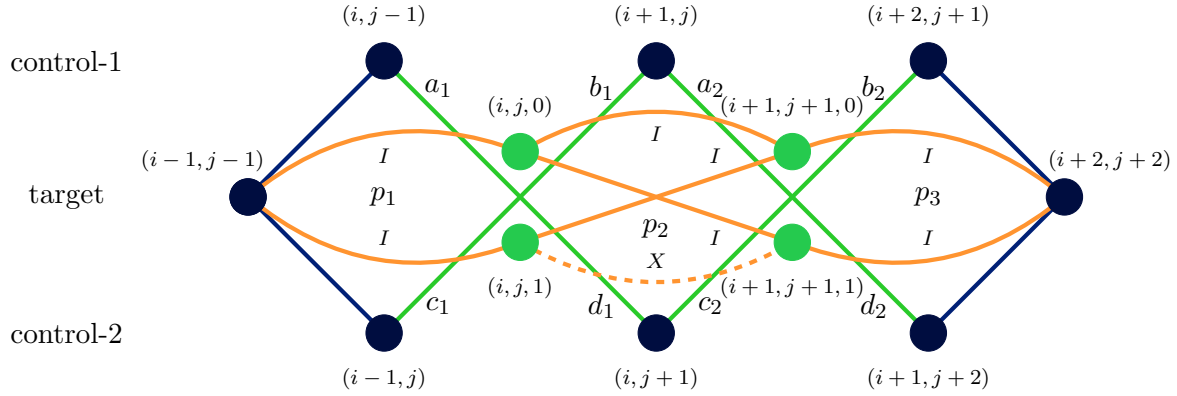


Figure 10: Direct Toffoli gate.

in Fig. 10. We consider the realization of a Toffoli gate in which the target qubit is *sandwiched* between the two control qubits. As seen in Fig. 10 in order to do the Toffoli gate we need to modify a larger region compared to the CNOT. Specifically, we now split two sites in two. In the first plaquette p_1 we apply the first step of the CNOT choosing the control-1 particle as control. There is one difference, namely we do not immediately apply the X -gate, but we just conditionally direct the target particle to site $(i, j, 0)$ or site $(i, j, 1)$ depending on the state of the control-1 particle. After this first intermediate step, we continue to do conditional operations depending on the state of the second control particle, which, in the meantime, needs to have completed a hopping to site $(i, j + 1)$ in order for the target particle to move forward. As an example, if the target particle is at $(i, j, 1)$ and the second control is the $s = 1$ state, the target particle will be allowed to only hop to the site $(i + 1, j + 1, 1)$ with the application of an X -gate, thus correctly realizing the logic of the Toffoli gate. The system works similarly in the remaining three cases in which the identity gate is applied. From the second split-sites $(i + 1, j + 1, 0)$ and $(i + 1, j + 1, 1)$ the target hops to the final single site at $(i + 2, j + 2)$ with the application of an identity gate. This construction of the Toffoli can be viewed as a double, sequential railroad switch.

In Appendix A we work these ideas out mathematically, showing that the effective Hamiltonian, obtained in lowest-order perturbation theory, applies the correct Toffoli gate logic. It is clear that with the same idea we can implement any controlled-controlled- U gate, just by substituting the hopping term that implements the X -gate with a hopping term that implements a generic single-qubit unitary as described in Sec. 2.

2.2. Dual Rail Representation

We represent a particle with spin $s = 0, 1$ by a pair of transmon qubits (qubit $s = 0$ or qubit $s = 1$) with an excitation in either one of the qubits. In this dual-rail representation, the hopping terms across a plaquette p implementing U_p , Eq. (4), become

simple flip-flop terms moving the excitation

$$\sum_{s,s'} \langle s' | U_p | s \rangle \sigma_{s'}^+[i+1, j+1] \sigma_s^-[i, j] + \text{h.c.}$$

This mapping leads to the couplers shown in Fig. 3 where U_p can be the X , I or H gate.

In the dual-rail representation, the total number operator $\mathbf{n}[\nu]$ for a particle at site ν is the sum of the number operators for the two transmon qubits at the site, counting whether a single excitation is present for the pair or none. The internal-state dependent number operator $n_s[\nu]$ for $s = 0, 1$ just counts whether transmon qubit $s = 0, 1$ has an excitation or not. This implies that the standard attractive edge terms can be written as doubled cross-Kerr couplers as in Fig. 4. For the CNOT, due to the split-site, one then requires several more cross-Kerr interactions as shown in Fig. 5. Realizing this connectivity is clearly a challenging aspect of our proposal.

If we use transmon qubits, we additionally have local $-\Omega \frac{\sigma_z}{2}$ terms in the Hamiltonian (see Eq. (17) and Eq. (23)) besides engineered flip-flop and cross-Kerr interactions. If all transmon qubits have identical frequencies, we can move to the rotating frame associated with that frequency without changing the Hamiltonian. This follows because any such rotating frame change leaves the cross-Kerr terms invariant while the on-track flip-flop interactions only depend on difference in frequencies between qubits on a track. Thus, in fact, it suffices to demand that all qubits on one track have the same frequency, in this case these extra single-qubit Z terms do not affect the dynamics.

As suggested in Table 1 and indicated with colors in Fig. 2 we imagine keeping qubits on adjacent tracks quite far detuned to avoid spurious intertrack excitation hopping.

Even the qubits on a single track can have small differences in frequencies: in a chosen rotating frame it implies that the hopping is slightly off-resonant and thus less effective. We numerically study this effect in Sec. 3.1.

3. Numerical Studies of Small Lattices

In this subsection we analyze errors in the model in Section 2, which originate from the fact that our effective Hamiltonian is arrived at in lowest-order perturbation theory. We start by looking at the dynamics in the rotated grid without the application of any gates. We study the probability that the string becomes disconnected since we intuitively expect it to relate to the probability of incorrect logic if we were to apply gates like the CNOT or Toffoli.

For small lattice sizes we have numerically studied the time-averaged probability for the string to be disconnected, defined as

$$\overline{P}_D(t) = \frac{1}{t} \int_0^t dt' P_D(t'). \quad (8)$$

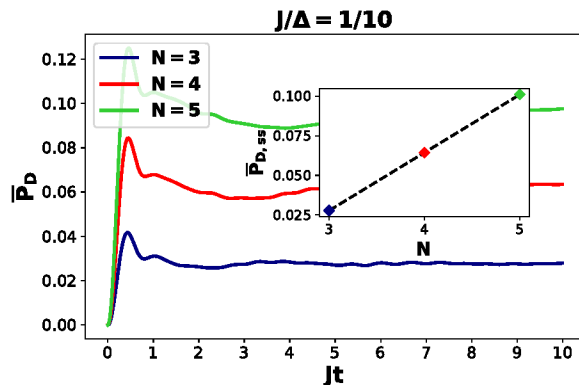


Figure 11: Time-averaged probability $\overline{P}_D$ for the string to be disconnected as a function of time for different lattice sizes. The inset shows the scaling of the steady state probability $\overline{P}_{D,ss}$, evaluated at a final time, with N .

We consider the time-averaged probability since $P_D(t)$ itself is an oscillatory function in time, demonstrating that these string-disconnect errors are coherent errors not accumulating in time. Fig. 11 shows this time-averaged probability $\overline{P}_D(t)$ as a function of time t for different lattice sizes, fixing $J/\Delta = 1/10$. We see that, unlike the instantaneous probability $P_D(t)$, the time-averaged probability reaches a steady value after a time of few $1/J$. In the inset we see that this steady-state average probability $\overline{P}_{D,ss}$ scales linearly with the lattice size for these lattice sizes.

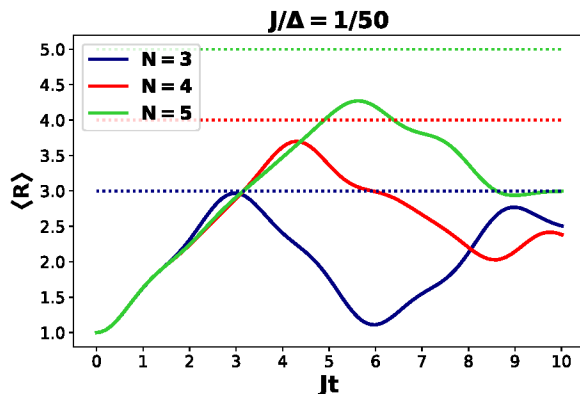


Figure 12: Average position of the particle on the central track as a function of time for different lattice sizes N . The dashed horizontal lines identify the maximum position that can be reached for each lattice size. $\langle R \rangle = 1$ is the initial position.

We further analyze the dynamics of the string of particles in Fig. 12. We plot the instantaneous average position of the particle on the central track, which is representative of the speed of the string. Formally $R = \sum_{i=1}^N \mathbf{n}[i, i]/N$. In this plot we take J/Δ to be quite small, i.e. $J/\Delta = 1/50$, so that the string is unlikely to be disconnected. We see that at short times of order $\sim 1/J$ the particle moves at a constant velocity, independent of lattice size, approximately equal to $0.6/(Jt)$ (lattice sites per

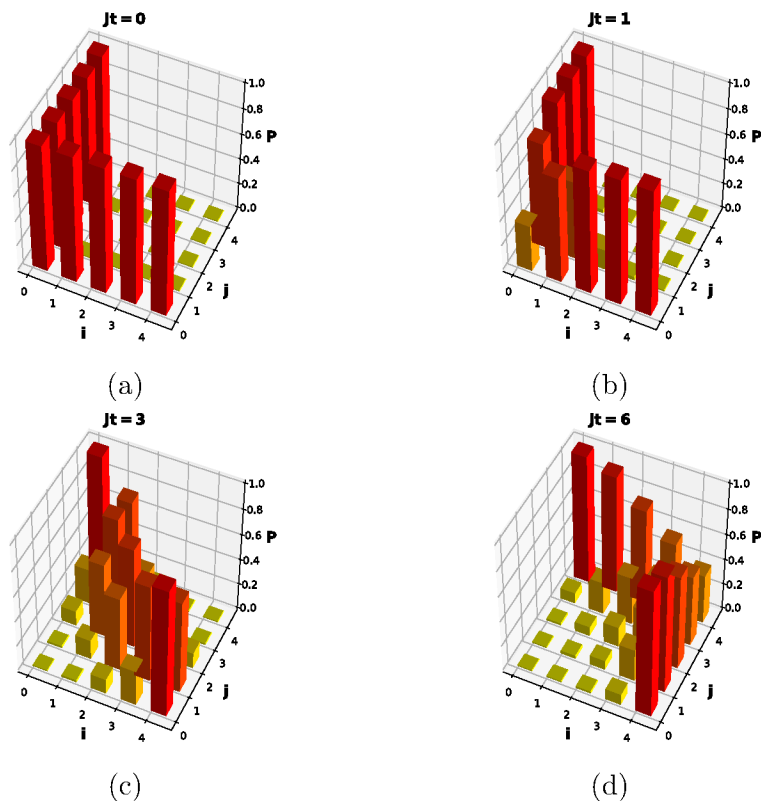


Figure 13: Plots of wavefront at different times. The bars represent the probability to find a particle in that position.

unit of dimensionless time Jt). This agrees with the analysis of Ref. [10] in which, based on the solution for $N \rightarrow +\infty$, it is argued that the string should move at constant velocity. In Fig. 18 we see how the average position of the particle on the middle track oscillates over time for much longer times (in red).

In Fig. 13 we further show some screenshots of the wavefront for lattice size $N = 5$ at different short times. We see that the particles tend to stay together as expected and move forward in a correlated way.

We now numerically examine the effect of perturbative errors on the execution of a CNOT gate as in Fig. 9. In our simulations we remove the spectator track at the bottom of Fig. 9. We then initialize the control and the target particles on the left sites with their spins in one fixed basis state. Our main goal is to evaluate the probability that the CNOT is implemented successfully or not. To this end we are interested in the probability that the particles are found on the right with the correct CNOT logic implemented on their internal states. We are interested in relatively short ‘first-passage’ times $\sim 1/J$ since we know that the string moves at a speed $\sim 1/J$.

We show a typical time evolution in Fig. 14, where we plot the probability of success P_S and the probability of error P_E defined as follows. P_E is the probability of finding both particles on the very right side, at the end of the CNOT region, but with wrong internal states according to the CNOT logic, while the success probability P_S is the

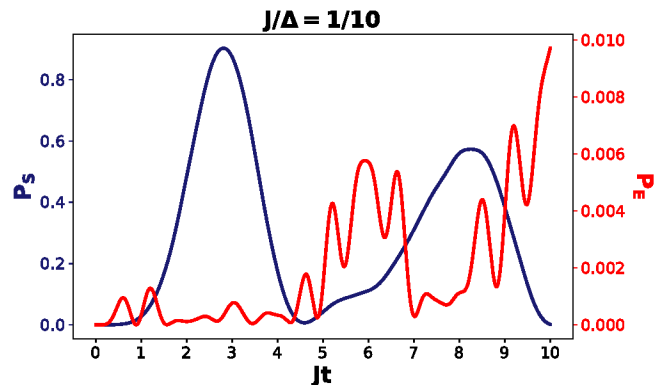


Figure 14: Example of probability of success (blue line, left axis) and error (red line, right axis) as a function of time for $J/\Delta = 1/10$.

probability for finding them on the right but in the correct state instead. Needless to say, the probability for arrival on the right $P_S + P_E \leq 1$. We see that the probability of success quickly increases, reaching a high maximum above at approximately $Jt = 3$, while the error stays relatively low, although we start to see an increase at the end of the simulation, reaching $\approx 1\%$.

In Fig. 15 we examine the scaling of the time-averaged error $\overline{P}_E(t)$ as a function of the ratio J/Δ . Naturally, we expect this error to decrease when we decrease the ratio J/Δ while looking at the same Jt . Fig. 15 shows that this probability scales as $(J/\Delta)^4$, and we find that this scaling is insensitive to our choice of Jt .

It is important to understand that the CNOT construction works under the assumption that the forward motion of the computation is guaranteed, so that the gate is not done and then undone many times. Since we only simulate a small CNOT region of a few sites, such accumulation of error by doing and undoing the gate incorrectly cannot be prevented, and is thus not representative of the expected behavior of the CNOT embedded in a large lattice. For example, Fig. 16 shows that when we average over very long times in our simulation, $\overline{P}_E$ and $\overline{P}_S$ become the same. This is a reflection of the fact that incorrect spin states and correct spin states at the sites on the right of the CNOT region have the same energy. It also shows that the probability that the CNOT gate fails becomes 1 in this long-time limit.

3.1. Static Disorder

We numerically examine how static disorder modifies the string dynamics. We focus on two kinds of disorder, namely variations in the hopping parameter J and variations of the on-site energy. In the dual-rail representation, the latter corresponds to variations of the transmon qubit frequency between qubits on different sites, possibly on one track. This model does not consider disorder in the qubit frequencies of a pair of transmon qubits, which should in principle be equal. Said in the language of the model in Sec. 2, we do not examine spin-dependent disorder. It is known that on-site or hopping disorder of the

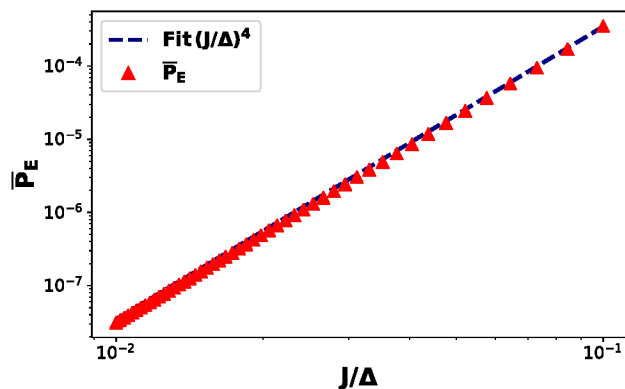


Figure 15: Time-averaged probability of error, as a function of J/Δ for CNOT at $Jt = 3$. We use a logarithmic scale on both axes.

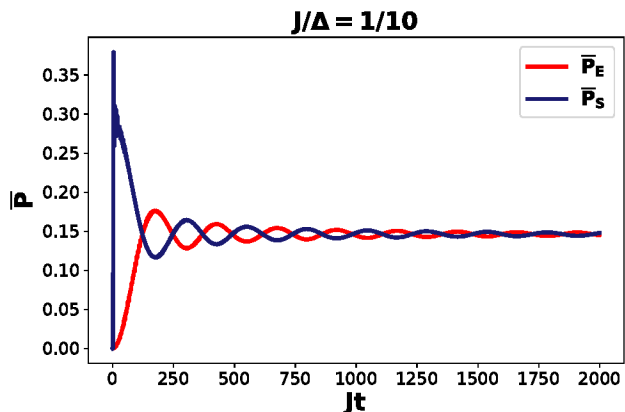
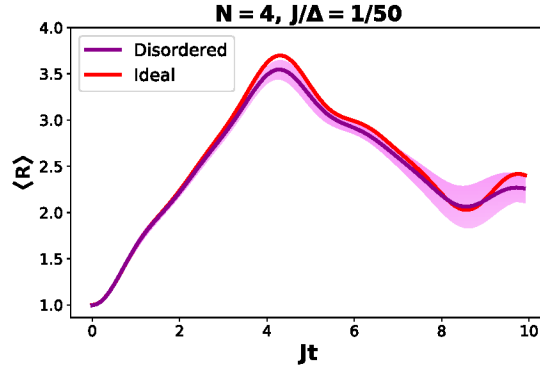


Figure 16: Long-time behaviour for the time-averaged probabilities of error and success.

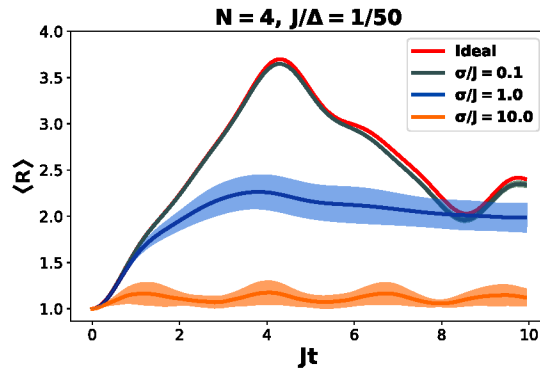
dynamics of a single particle on a line leads to Anderson localization with a localization length which depends on the disorder strength. Here the question is to understand what localization of the string occurs at what disorder strength and when this localization would prevent the execution of the computation.

In Fig. 17 we show results for the case $N = 4$: we compare the ideal time evolution of the average position of the particle on the central track with the disordered one. For disorder in the hopping parameter, we take J to vary by roughly 10%, which is reasonable in our transmon implementation. In Fig 17a we see that this amount of disorder slightly influences the dynamics of the string, but it still reaches the center of the grid with approximately the same velocity as in the ideal case.

The situation is different in the case of disorder on the on-site energy as shown in Fig. 17b. As argued previously, when all on-site energies, translating into transmon qubit frequencies, are identical, one can rotate away this frame. When frequencies differ, the flip-flop coupling will be off-resonant and thus less effective in bringing about forward motion. In Fig. 17b we see that increasing the ratio between the standard deviation and the hopping parameter, we pass from a configuration in which there is essentially



(a) Ideal time evolution vs. average time evolution with 10% Gaussian disorder on the hopping parameter J (meaning that the standard deviation of the Gaussian is taken to be $\sigma = J \times 0.1$).



(b) Ideal time evolution vs. time evolution with Gaussian disorder on the on-site energy. For $\sigma/J = 0.1$ the light shaded area is not visible by eye.

Figure 17: Ideal vs. disordered time evolution. The disordered data are averaged over 50 simulation runs. The light shaded areas represent the standard deviation of the disordered data at each time, which shall not be confused with the standard deviation σ of the distribution of J . $\langle R \rangle = 1$ is the initial position.

no localization ($\sigma/J = 0.1$) to a localized configuration ($\sigma/J = 10.0$). This tells us that in order for the string to propagate we need an on-track variation of the on-site energy that is $\leq J$. In our transmon implementation and parameters chosen as in Table 1, this means a variation of no more than 15 MHz on the qubit frequencies, which can be achieved by careful design.

Finally, Fig. 18 shows the long-time behaviour of the position of the central particle for the ideal case versus the case with disorder of the on-site energy. In the ideal case, while on average the particle is in the middle, it still presents oscillations inherent to the

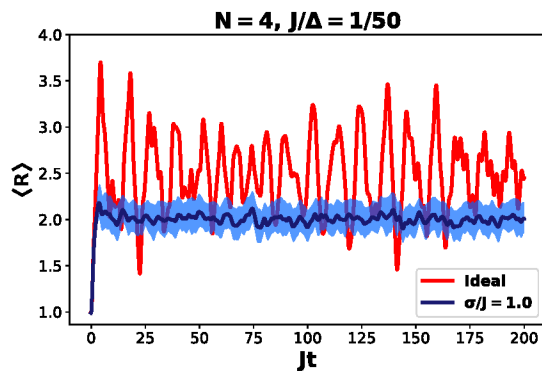


Figure 18: Ideal long-time evolution vs. on-site disordered long-time evolution. As in Fig. 17, the light shaded area represents the standard deviation of the disordered data at each time.

unitary dynamics. The averaged disorder evolution instead presents little oscillations and we also notice that average position is not exactly in the middle (the middle position is $\langle R \rangle = 2.5$), but slightly closer to the initial position.

From these numerics we cannot readily extrapolate what happens for larger N . The interesting effect of small disorder is to localize the string where indeed more localization occurs with more disorder. An open question is how the localized $\langle R \rangle$ depends on N , i.e. does the stationary value for $\langle R \rangle$ at fixed disorder strength decrease as a function of N (and if so, as what function).

We have not included numerical studies of disorder on the cross-Kerr coupling. As long as Δ/J is sufficiently large to warrant the perturbative picture, we do not expect such disorder to play a large role in the string dynamics.

4. Cross-Kerr and flip-flop couplings between superconducting transmon qubits

In this section we describe details of the required couplers to implement the Hamiltonian quantum computing scheme using superconducting transmon qubits.

4.1. Cross-Kerr interaction

The most challenging interaction to engineer in our problem is the strong cross-Kerr interaction between transmon qubits. There are several proposals for the implementation of this kind of interaction with transmon qubits in the literature [29, 30, 31, 32, 33, 19, 34]. In particular, we will build on ideas similar to Refs. [31, 32, 33, 19], in which a simple junction in parallel with a capacitance is used to implement the desired cross-Kerr interaction. This element also induces a linear flip-flop interaction which is undesired. However, the capacitance and the Josephson junction give rise to flip-flop terms of opposite signs and consequently we can find a value of the

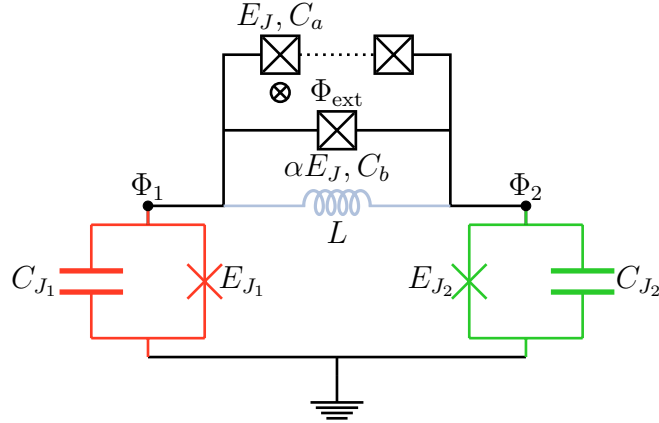


Figure 19: Attractive cross-Kerr coupler between two grounded transmon qubits. The parameter $\alpha < 1$ is the ratio of the (small) black-sheep junction Josephson energy versus the Josephson energy of a junction in the array. The different colors of the transmons denote that they should have different frequencies as shown in the color scheme in Fig. 2. Additionally, the total coupling capacitance is $C_c = C_a/N_J + C_b$. The inductance is shown in grey as it is not an essential element of the coupler and could be omitted. Φ_1 and Φ_2 represent the node flux variables, while Φ_{ext} the external flux in the loop formed by the junction and the array of junctions.

capacitance that exactly cancels the linear term. This is basically the same idea of an LC filter, which is based on the fact that the impedance of a capacitor and an inductor have different signs [35].

In order to limit the value of the capacitance needed to achieve this cancellation, and thus the cross-talk between the qubits on the same track (see discussion in Subsec. 4.2), we consider, instead of a simple junction, a modified direct coupler shown in Fig. 19. It represents two transmons coupled by a junction in parallel with an array of N_J junctions and an inductance L .

The transmons have different frequencies since they sit on adjacent tracks according to our scheme in Fig. 2. In the following paragraphs we show how to obtain the Hamiltonian of two coupled transmon qubits from the circuit in Fig. 19 through various approximations: the final Hamiltonian is in Eq. (23).

We assume that the array of junctions is operated in the limit of $E_J/E_C \gg 1$, with $E_{C_a} = e^2/(2C_a)$, as well as the limit in which the resonant frequencies of the array are larger than any other frequency of the problem. In this limit, the internal degrees of freedom of the array can be eliminated and the array can be treated with an effective potential

$$U_{\text{array},m}(\varphi) = -N_J E_J \cos\left(\frac{\varphi + 2\pi m}{N_J}\right), \quad (9)$$

with φ the superconducting phase difference across the element and $m \in \{0, 1, \dots, N_J - 1\}$ [36, 37, 38, 39]. The effective Hamiltonian depends on the parameter m , which labels the different classical metastable minima of the total potential of the array in the limit

of $E_J/E_{Ca} \rightarrow +\infty$. In the phase-slip model of the array of junctions, first discussed in Ref. [40], we associate a quantum state $|m\rangle$ with energy given by Eq. (9) to each metastable minimum of the array. The index m of a minimum represents the number of vortices, i.e., the number of 2π -turns that the phase along the array undergoes.

In the limit of $E_J/E_{Ca} \gg 1$ (classical limit), we expect that given a certain state $|m\rangle$ the amplitude of tunneling to a different $|m'\rangle$ is small and we can effectively assume that the potential is given by Eq. (9). This is exactly the same working regime of the superinductances used for the fluxonium qubit [41]. In what follows we will assume that the array of junctions is initially set in the state $|m=0\rangle$ ||, so that our effective array potential is (see also Ref. [17])

$$U_{\text{array},0}(\varphi) = -N_J E_J \cos \frac{\varphi}{N_J}$$

The Josephson junction array will also add an additional capacitance given by C_a/N_J , with C_a the capacitance of a single junction in the array. The total coupling capacitance is then $C_c = C_a/N_J + C_b$, with C_b the capacitance of the small (black sheep) junction in parallel. It is worth mentioning that a system composed of an array of three large junctions in parallel with a black sheep junction has been analyzed in Refs. [18, 42], and nicknamed the SNAIL, with the goal of obtaining a potential that gives rise to a three-wave mixing term (third-order in annihilation/creation operators), but without cross-Kerr (quartic). Here instead we would like to limit the quadratic term, while keeping the cross-Kerr interaction.

We obtain the Lagrangian of the circuit in Fig. 19 as

$$\mathcal{L} = \frac{C_{J_1}}{2} \dot{\Phi}_1^2 + \frac{C_{J_2}}{2} \dot{\Phi}_2^2 + \frac{C_c}{2} (\dot{\Phi}_1 - \dot{\Phi}_2)^2 - U_{\text{tot}}(\Phi_1, \Phi_2),$$

with the total potential

$$U_{\text{tot}}(\Phi_1, \Phi_2) = -E_{J_1} \cos \left[\frac{2\pi}{\Phi_0} \Phi_1 \right] - E_{J_2} \cos \left[\frac{2\pi}{\Phi_0} \Phi_2 \right] + \frac{1}{2L} \left(\Phi_1 - \Phi_2 \right)^2 - \alpha E_J \cos \left[\frac{2\pi}{\Phi_0} (\Phi_1 - \Phi_2) \right] - N_J E_J \cos \left[\frac{2\pi}{\Phi_0} \frac{\Phi_1 - \Phi_2 + \Phi_{\text{ext}}}{N_J} \right]. \quad (10)$$

We define the conjugate variables $Q_{1,2} = \frac{\partial \mathcal{L}}{\partial \dot{\Phi}_{1,2}}$, in terms of which the Hamiltonian (obtained via a Legendre transform) reads

$$H = \frac{Q_1^2}{2\tilde{C}_{J_1}} + \frac{Q_2^2}{2\tilde{C}_{J_2}} + \frac{Q_1 Q_2}{\tilde{C}_c} + U_{\text{tot}}(\Phi_1, \Phi_2),$$

where we defined the capacitances

$$\frac{1}{\tilde{C}_{J_1}} = \frac{C_{J_2} + C_c}{\det(\mathbf{C})}, \quad \frac{1}{\tilde{C}_{J_2}} = \frac{C_{J_1} + C_c}{\det(\mathbf{C})}, \quad \frac{1}{\tilde{C}_c} = \frac{C_c}{\det(\mathbf{C})},$$

|| Note that we could also start with a different $|m\rangle$ and obtain the same results, but with external fluxes shifted by an integer of 2π .

with the determinant of the capacitance matrix $\mathbf{C}$ as

$$\det(\mathbf{C}) = C_{J_1}C_{J_2} + (C_{J_1} + C_{J_2})C_c.$$

We rewrite $Q_{1,2} = 2e\mathbf{n}_{1,2}$ with $\mathbf{n}_m$ the number of Cooper pairs and the superconducting phases $\varphi_{1,2} = 2\pi\Phi_{1,2}/\Phi_0$. In order to quantize the problem we promote $\mathbf{n}_m$ and φ_m to operators with commutation relation $[\hat{\varphi}_m, \hat{\mathbf{n}}_m] = iI$, but we will continue to write $\hat{\mathbf{n}}$ as $\mathbf{n}$ and $\hat{\varphi}$ as φ .

Using this notation, we rewrite our Hamiltonian as

$$H = 4E_{C_1}\mathbf{n}_1^2 + 4E_{C_2}\mathbf{n}_2^2 + 8E_{\text{cap}}^{\text{coup}}\mathbf{n}_1\mathbf{n}_2 + U_{\text{tot}}(\varphi_1, \varphi_2), \quad (11)$$

with charging energy $E_{C_m} = e^2/(2\tilde{C}_{J_m})$ and a capacitive coupling energy between the two transmons equal to

$$E_{\text{cap}}^{\text{coup}} = e^2/(2\tilde{C}_c). \quad (12)$$

Let us now focus on the coupling part of the potential U_{tot} in terms of the phase difference $\varphi = \varphi_1 - \varphi_2$ at the nodes in Fig. 19:

$$U_c(\varphi) = \frac{E_L}{2}\varphi^2 - \alpha E_J \cos \varphi - N_J E_J \cos\left(\frac{\varphi + \varphi_{\text{ext}}}{N_J}\right),$$

Here we introduced the inductive energy $E_L = \Phi_0^2/(4\pi^2 L)$. Additionally, we fix the external flux to the value $\varphi_{\text{ext}} = 2\pi\Phi_{\text{ext}}/\Phi_0 = N_J\pi$. We will now assume that it is possible to Taylor-expand the coupling potential up to fourth order, as it is usually done for transmon qubits. This is a good approximation as long as we work in the transmon regime $E_{J_m}/E_{C_m} \gg 1$. In addition, we should also guarantee that the total potential has a global minimum at the origin. Expanding the coupling potential around $\varphi = 0$ gives

$$U_c(\varphi)/E_J = -\alpha + N_J + \frac{1}{2}\left[\alpha + \frac{E_L}{E_J} - \frac{1}{N_J}\right]\varphi^2 - \frac{1}{24}\left[\alpha - \frac{1}{N_J^3}\right]\varphi^4 + \mathcal{O}(\varphi^6).$$

We see that by setting

$$\alpha + \frac{E_L}{E_J} - \frac{1}{N_J} = 0. \quad (13)$$

the quadratic term vanishes completely, while the quartic term would be approximately the same as the case with a simple junction, since the contribution due to the array scales as $1/N_J^3$. One should keep in mind that this point is a maximum of the coupling potential, but always (at least) a relative minimum of the total potential. In fact, the reason for including the inductor L in the circuit of Fig. 19 is to make sure that the origin is a global minimum of the total potential, and not only a metastable minimum. In Fig. 20 we plot an example of this total potential showing that it has a global minimum in the origin and this holds true for all the parameters that we have considered in the next subsection. However, the inductance is not necessary if we accept to work in a metastable minimum.

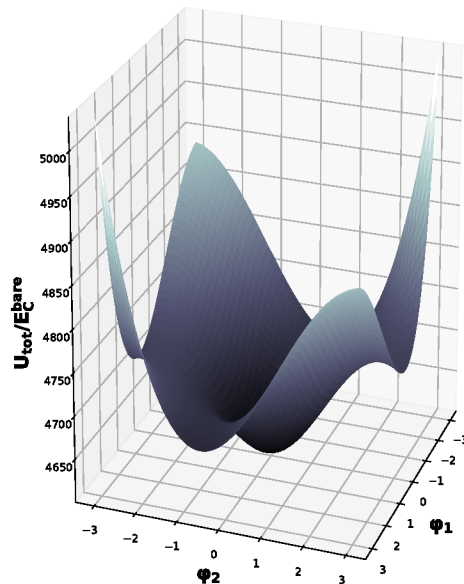


Figure 20: Total potential $U_{\text{tot}}(\varphi_1, \varphi_2)$ for the parameters in Table 2 where α is set to 0.043, showing that the global minimum occurs at the origin.

As we will see, we shall aim to meet the condition in Eq. (13) only approximately as the nonlinearity φ^4 renormalizes the hopping strength in the transmon qubit subspace. It is convenient to define the following energies

$$E_{\text{ind}}^{\text{coup}} = E_J \left(\alpha + \frac{E_L}{E_J} - \frac{1}{N_J} \right), \quad (14a)$$

$$E_{\text{Kerr}}^{\text{coup}} = E_J \left(\alpha - \frac{1}{N_J^3} \right). \quad (14b)$$

We now perform a Taylor expansion for the total Hamiltonian in Eq. 11 and rewrite it as

$$H = \sum_{m=1}^2 H_m + H_{\text{lin}} + H_{CK} + H_{\text{non.lin}},$$

where we defined the Hamiltonian of the single transmons as

$$H_m = 4E_{C_m} \mathbf{n}_m + \frac{E_{L_m}}{2} \varphi_m^2 - \frac{E_{K_m}}{24} \varphi_m^4,$$

with for $m = 1, 2$

$$E_{L_m} = E_{J_m} + E_{\text{ind}}^{\text{coup}}, E_{K_m} = E_{J_m} + E_{\text{Kerr}}^{\text{coup}}. \quad (15)$$

The linear coupling Hamiltonian H_{lin} is given by

$$H_{\text{lin}} = 8E_{\text{cap}}^{\text{coup}} \mathbf{n}_1 \mathbf{n}_2 - E_{\text{ind}}^{\text{coup}} \varphi_1 \varphi_2,$$

while the term responsible for the cross-Kerr interaction is

$$H_{\text{CK}} = -\frac{E_{\text{Kerr}}^{\text{coup}}}{4} \varphi_1^2 \varphi_2^2.$$

We also introduced what we called a nonlinear Hamiltonian

$$H_{\text{non.lin}} = \frac{E_{\text{Kerr}}^{\text{coup}}}{6} (\varphi_1^3 \varphi_2 + \varphi_1 \varphi_2^3).$$

We introduce annihilation and creation operators for the transmon modes

$$\begin{aligned} \varphi_m &= \left(\frac{2E_{C_m}}{E_{L_m}} \right)^{1/4} (a_m + a_m^\dagger), \\ \mathbf{n}_m &= \frac{i}{2} \left(\frac{E_{L_m}}{2E_{C_m}} \right)^{1/4} (a_m^\dagger - a_m). \end{aligned}$$

We now rewrite the previously defined Hamiltonian using annihilation and creation operators, performing additionally several rotating wave approximations (RWA):

- Transmon Hamiltonian:

$$\frac{H_m}{\hbar} \stackrel{\text{RWA}}{=} (\omega_m + \delta_m) a_m^\dagger a_m + \frac{\delta_m}{2} a_m^\dagger a_m^\dagger a_m a_m, \quad (17)$$

with the transmon frequency $\omega_m = \sqrt{8E_{C_m}E_{L_m}}/\hbar$ and the anharmonicity $\delta_m = -E_{K_m}E_{C_m}/(\hbar E_{L_m})$.

- Linear coupling Hamiltonian:

$$\frac{H_{\text{lin}}}{\hbar} \stackrel{\text{RWA}}{=} J_{\text{lin}} (a_1 a_2^\dagger + a_1^\dagger a_2),$$

with the linear hopping parameter J_{lin} given by

$$J_{\text{lin}} = J_{\text{cap}} + J_{\text{ind}}, \quad (18)$$

where we have a capacitive and an inductive J_{ind} contribution, i.e.

$$J_{\text{cap}} = \frac{1}{\hbar} 2E_{\text{cap}}^{\text{coup}} \left(\frac{E_{L_1}}{2E_{C_1}} \right)^{1/4} \left(\frac{E_{L_2}}{2E_{C_2}} \right)^{1/4}, \quad (19a)$$

$$J_{\text{ind}} = -\frac{1}{\hbar} E_{\text{ind}}^{\text{coup}} \left(\frac{2E_{C_1}}{E_{L_1}} \right)^{1/4} \left(\frac{2E_{C_2}}{E_{L_2}} \right)^{1/4}. \quad (19b)$$

In these equations we clearly see that capacitive and inductive coupling give a flip-flop interaction of opposite sign.

- Cross-Kerr Hamiltonian

$$\frac{H_{\text{CK}}}{\hbar} \stackrel{\text{RWA}}{=} J_{\text{CK}} \left(a_1^\dagger a_1 a_2^\dagger a_2 + \frac{1}{2} a_1^\dagger a_1 + \frac{1}{2} a_2^\dagger a_2 + \frac{1}{4} a_1 a_1 a_2^\dagger a_2^\dagger + \frac{1}{4} a_1^\dagger a_1^\dagger a_2 a_2 \right), \quad (20)$$

with

$$J_{\text{CK}} = -\frac{1}{\hbar} E_{\text{Kerr}}^{\text{coup}} \left(\frac{2E_{C_1}}{E_{L_1}} \right)^{1/2} \left(\frac{2E_{C_2}}{E_{L_2}} \right)^{1/2} \equiv -\Delta. \quad (21)$$

In Eq. (20) the RWA means that we have omitted all fast rotating terms which are products of unequal numbers of creation and annihilation operators. The last two terms in Eq. (20) which represent off-resonant two-photon processes contribute little and disappear when projecting onto the transmon qubit subspaces. We see that when the transmon qubits become more harmonic, i.e. E_J/E_{C_m} increases for $m = 1, 2$, J_{CK} becomes progressively smaller as E_{L_m}/E_{C_m} defined through Eq. (15) increases.

- Nonlinear Hamiltonian

$$\begin{aligned} \frac{H_{\text{non.lin}}}{\hbar} \stackrel{\text{RWA}}{=} J_{\text{non.lin}}^{(1)} [(a_1 a_2^\dagger + a_1^\dagger a_2) + (a_1 a_2^\dagger + a_1^\dagger a_2) a_1^\dagger a_1 + \\ a_1^\dagger a_1 (a_1 a_2^\dagger + a_1^\dagger a_2) + (a_1 a_2^\dagger a_1^\dagger a_1 + a_1^\dagger a_1 a_1^\dagger a_2)] + \\ J_{\text{non.lin}}^{(2)} [(a_1 a_2^\dagger + a_1^\dagger a_2) + (a_1 a_2^\dagger + a_1^\dagger a_2) a_2^\dagger a_2 + \\ a_2^\dagger a_2 (a_1 a_2^\dagger + a_1^\dagger a_2) + (a_1 a_2^\dagger a_2^\dagger a_2 + a_2^\dagger a_2 a_1^\dagger a_2)], \end{aligned}$$

with for $m = 1, 2$

$$\begin{aligned} J_{\text{non.lin}}^{(1)} &= \frac{1}{\hbar} \frac{E_{\text{CK}}}{6} \left(\frac{2E_{C_1}}{E_{L_1}} \right)^{3/4} \left(\frac{2E_{C_2}}{E_{L_2}} \right)^{1/4} \\ J_{\text{non.lin}}^{(2)} &= \frac{1}{\hbar} \frac{E_{\text{CK}}}{6} \left(\frac{2E_{C_1}}{E_{L_1}} \right)^{1/4} \left(\frac{2E_{C_2}}{E_{L_2}} \right)^{3/4}. \end{aligned}$$

Again RWA means that fast-rotating terms with unequal numbers of creation and annihilation operators are omitted. In the transmon qubit space all these nonlinear terms act as flip-flop interactions.

When we project the various terms of the full Hamiltonian onto the first two levels of each transmon qubit, we obtain the final Hamiltonian

$$\frac{H_Q}{\hbar} = -\frac{\Omega_1}{2} \sigma_1^z - \frac{\Omega_2}{2} \sigma_2^z + J_{\text{hop}} (\sigma_1^+ \sigma_2^- + \sigma_1^- \sigma_2^+) - \Delta n_1 n_2, \quad (23)$$

Here the fully-dressed transmon frequency is $\Omega_m = \omega_m + \delta_m - \Delta/2$, while the total *forbidden* hopping strength is

$$J_{\text{hop}} = J_{\text{cap}} + J_{\text{ind}} + J_{\text{non.lin.}}, \quad (24)$$

with

$$J_{\text{non.lin}} = 3 \sum_m J_{\text{non.lin}}^{(m)}. \quad (25)$$

Since we never drive the transmon qubits, except for creating the initial excitations, the qubit approximation is justified.

We observe that the dressed transmon frequency will depend on the number of cross-Kerr interactions that a transmon qubit participates in. This means that if we want transmon qubits to have the same dressed frequencies, then the bare transmon frequency $\omega_m + \delta_m$ in Eq. (17) needs to be different if the qubit participates in a different number of cross-Kerr couplers. Qubits which differ in this way occur at the boundaries of the computational lattice in Fig. 2 as well as on the split-sites in the CNOT (or Toffoli) region shown in Fig. 5.

4.2. Cross-Kerr Coupling Strength and Cross-Talk

When we consider our coupler in a larger system we have to analyze the problem of unwanted long-range coupling (cross-talk), which is always potentially present in superconducting qubit architectures. In order to mitigate this problem one requires that the coupling capacitances C_c are much smaller than the transmon capacitances C_J , so that the inversion of the capacitance matrix remains local in a first-order approximation in C_c/C_J .

Cross-talk via two cross-Kerr couplers would be particularly undesirable for qubits on the same track as these have the same frequency, see e.g. Fig. 21. In order to limit this problem, our idea is to limit the inductive contribution to the linear coupling while keeping the cross-Kerr approximately the same, which is what the coupler analyzed in the previous subsection achieves. This should limit the capacitance needed to *filter* the hopping interaction between different tracks and thus moderate the problem of cross-talk.

There are however several trade-offs that one should consider in choosing parameter strengths. First of all, we need to require the ratio E_J/E_{C_a} of each junction of the array to be much larger than one, in order to prevent phase slips. We will always fix $E_J/E_{C_a} = 100$. This is in contrast with our desire to have a small effective coupling capacitance. To reduce this capacitance we could also increase the number of junctions, but from a practical point of view we would like to limit N_J . In what follows, we just assume a set of typical and reasonable experimentally achievable parameters, and compute the typical forbidden hopping strength, the cross-Kerr strength as well as the cross-talk hopping induced by two couplers. In particular, in order to evaluate the order of magnitude of this cross-talk hopping, we consider a circuit as in Fig. 21 where two transmons on the same track (red) are both coupled to a transmon on a different track (green).

We will focus on the case $N_J = 4$ and take the bare charging energy of the bare transmons as in Fig. 19 to be equal, i.e. $E_C^{\text{bare}} \equiv e^2/(2C_J)$ as our unit of energy and also set $\hbar = 1$. We neglect the intrinsic capacitance of the coupling junction $C_b = 0$.

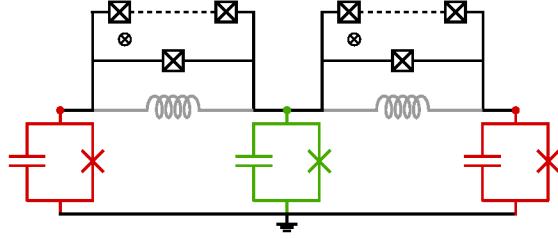


Figure 21: Two resonant transmons (red) coupled to a common detuned transmon (green). Parameters are not shown for simplicity. The resonant transmons can either be two transmons forming a pair or two transmon qubits on adjacent sites on a track.

E_{J_1}	80
E_{J_2}	60
E_{C_a}	12
E_J	1200
E_L	255

Table 2: Parameters in units of E_C^{bare} .

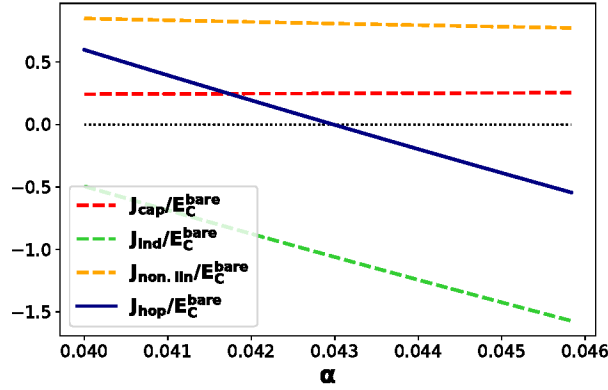
The chosen parameters are listed in Table 2. In Fig. 22a we see how it is possible to achieve a hopping term that is exactly equal to zero. This is possible with an effective capacitance that is much smaller than the transmon capacitance, in particular $C_c \approx 1/(48)C_J$. With these parameters we can get a high cross-Kerr interaction $\Delta/E_C^{\text{bare}} \approx 0.8$ as we can see from Fig. 22b. Assuming a typical bare charging energy $E_C^{\text{bare}}/h \approx 200\text{MHz}$, we then get $\Delta \approx 2\pi \times 160\text{MHz}$.

Considering the circuit in Fig. 21, and using again the parameters listed in Table 2, we estimate a cross-talk hopping between the red transmons of approximately $0.004E_C^{\text{bare}}$, which corresponds approximately to $2\pi \times 0.8\text{MHz}$. If these two red transmons form a pair, then this hopping should ideally be of zero strength. If these two red transmons are two transmons on the same track at nearest sites, then this term will weakly contribute to the total flip-flop interaction.

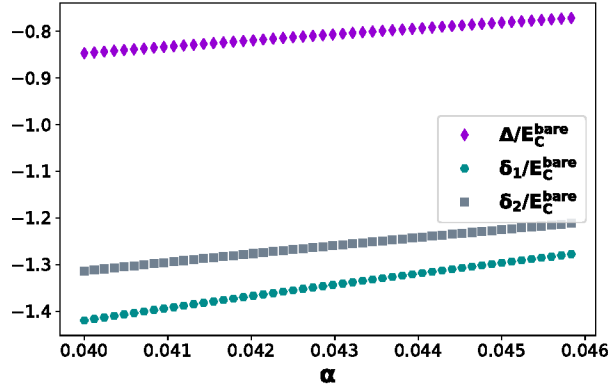
More precisely, if we choose the perturbative regime $J/\Delta = 1/10$, where J is now the desired on-track hopping strength of Table 1, the weak on-track flip-flop interaction that we discuss in the next section can be chosen as $J \approx 0.08E_C^{\text{bare}}$. Thus the contribution from the cross-talk ‘cross-Kerr mediated’ hopping is negligibly small.

4.3. Weak On-Track Flip-Flop Interactions

As we have seen in the previous section the hopping interaction is naturally obtained with transmon qubits, by simply using capacitances or inductances. As we know, we need hopping interactions between qubits on a track, which are much weaker than the cross-Kerr interactions on the edges. On the other hand, these interactions should be larger than the forbidden hopping and cross-talk hopping strengths.



(a) Different contributions to the forbidden hopping parameter with their strengths as defined in Eqs. 19 and Eq. 25 (dashed lines) and the total forbidden hopping strength (solid line) as defined in Eq. (24).



(b) Cross-Kerr coupling and anharmonicities of the two transmons (see Eqs. 17 and Eq. 21 for the definitions).

Figure 22: Forbidden hopping (a) and cross-Kerr (b) coupling strengths for parameters α close to the condition of canceling hopping in Eq. (13). In Fig. 22a we see that the hopping coupling is zero at $\alpha \approx 0.043$, corresponding to a Josephson energy of the small junction equal to $\alpha E_J/E_C^{\text{bare}} = 51.72$. At this point, $\Delta/E_C^{\text{bare}} \approx 0.8$.

The general form of the capacitive coupling between two equal transmons can be deduced from the analysis done in Sec. 4.1 and is of the general form (see Eq. (18)):

$$J_{\text{cap}} = \frac{1}{\hbar} 2E_{\text{cap}}^{\text{coup}} \left(\frac{E_L}{2E_C} \right)^{1/2}.$$

with $E_{\text{cap}}^{\text{coup}}$ in Eq. (12). In this expression we have assumed that the transmons are coupled via the cross-Kerr coupler to other parts of the circuits so that the inductive energy E_L includes the inductance due to this coupling as in Eq. (15). Similarly, one

can imagine that E_C is a dressed charging energy which includes the capacitive energy contributions from different couplings. Assuming $E_C/h \approx 200\text{MHz}$, we obtain the magnitude of a typical flip-flop interaction to be $J \approx 2\pi \times 10 - 15\text{MHz}$, which can be easily obtained with a small capacitance.

We see that a capacitive coupling always gives rise to a coupling of the same sign. We do however need a sign change to implement the Hadamard gate. An immediate solution would be to use inductances to implement couplings of negative signs. However, in order to keep the coupling small, one would need to work in the regime of large inductances and thus use arrays of junctions.

There is a much simpler way to achieve this sign change which is to couple both qubits capacitively to a common resonator, and require the qubits to be far detuned from the resonators' frequencies, working in the so-called dispersive regime. Let g_{res} be the coupling to the resonator of the two qubits, let $\delta < 0$ be their anharmonicity, let Ω be their frequency (which should be the same for qubits on one track) and let ω_{res} be the frequency of the resonator. One has an effective flip-flop coupling between the two qubits [43, 16] of strength

$$J_{\text{eff}} = \frac{g_{\text{res}}^2}{\Omega - \omega_{\text{res}}} - \frac{1}{2} \frac{(\sqrt{2}g_{\text{res}})^2}{\Omega + \delta - \omega_{\text{res}}},$$

It is clear that we can obtain couplings of different signs by properly selecting the frequency of the resonator ω_{res} . For instance by taking $\omega_{\text{res}} > \Omega$, we would get a negative effective hopping as expected. The circuit that implements the Hadamard gate is represented in Fig. 3. It is clear, that besides a Hadamard gate, any real rotation

$$U(\theta) = \begin{pmatrix} \cos(\theta) & \sin(\theta) \\ -\sin(\theta) & \cos(\theta) \end{pmatrix}$$

could be engineered using (resonator-mediated) flip-flop couplings. The two pairs of transmon qubits need to be four-way coupled as in Fig. 3b where the strength of the flip-flop interactions should be set to $J_{00} = J \cos(\theta)$, $J_{10} = J \sin(\theta)$ etc. where J is the global hopping strength.

4.4. Real Interactions and Time-Reversal

Since we obtain the flip-flop interaction from a passive capacitive coupling, i.e., we do not use any external drives, and we work in the transmon regime, we are giving up the possibility to have complex-valued couplings in our chosen transmon basis. This is a consequence of the fact that the Hamiltonians are invariant under a time-reversal transformation, and additionally, the basis we have chosen is the eigenbasis of an operator invariant under time-reversal symmetry, namely the basis of the uncoupled transmon Hamiltonians. More precisely, the capacitive coupling $Q_1 Q_2$ and the inductive coupling $\Phi_1 \Phi_2$ are invariant under the time-reversal transformations $Q_m \rightarrow Q_m$,

$\Phi_m \rightarrow -\Phi_m$ [44]. We may point out that passivity alone, meaning no external time-dependent drive, does not automatically imply reciprocity in circuit-QED systems as shown in Refs. [45, 46].

5. Challenges

As the present computational model does not specify a means for error correction, the effect of coherent errors and stochastic noise has immediate consequences for its viability. In the next paragraphs we discuss two specific challenges for using the proposed architecture for some form of quantum computation or quantum learning.

5.1. Loss of Excitations

As mentioned at the beginning of the paper and in [11], the loss of transmon excitations puts a severe limitation on the computation time. Suppose that we have a single hopping particle on a line with L sites. We associate with each site a decay rate γ that destroys the particle, i.e., it brings the system from one of the sites $|k\rangle$ to a vacuum state $|\text{vac}\rangle$. Formally, we can model such dissipation using a Lindblad master equation

$$\frac{d\rho}{dt} = -i[H_L, \rho] + \sum_{k=0}^{L-1} \gamma D[|\text{vac}\rangle \langle k|] \rho,$$

with the Lindblad dissipator $D[A]\rho = A^\dagger \rho A - 1/2\{A^\dagger A, \rho\}$ and H_L the Hamiltonian of a quantum walk on a line, see e.g. Eq. (B.1). It is not hard to show that this implies that the decay rate of the excitation is γ (since the single particle can only be annihilated at one of the sites), so that the decay rate does not depend on the length of the line L or the depth of the computation. Of course, the total probability of particle loss does still increase with the depth of the computation as the computation time linearly increases with the computational depth.

5.2. Arrival Time and Measurement

One question in the Hamiltonian computing models, both the single-clock as well as the multi-clock model, is to determine which qubits to measure to read out the computation. Ideally, the computation would finish in a known amount of time and only the relevant computational qubits are measured. For example, for the Feynman model as realized on the 2D lattice in Appendix B, these could be the qubits in the final column of the computation. However, it is known that, at a fixed time, the probability to get to the last site of the walk on a line with L sites decreases with $1/L$ and various ideas have been formulated to address this problem, see e.g. [47].

One simple method is to pad the quantum circuit with I gates, i.e. make the walk, say, twice as long, and measure at a random, sufficiently long, time to determine where the excitations reside. In this way, the probability for determining the total output of

the computation at a random time can be made arbitrarily close to 1. In the model in Appendix B it would correspond to measuring a whole 2D rectangular region of qubits at the end of the lattice. We call such region of qubits to be measured at a specific (random) time the measurement region.

Similarly, in [10] it was envisioned that the actual quantum circuit with non-trivial gates is encoded in a $K \times K$ region of the lattice located in the left corner with $K = N/4$, while in the rest of the lattice only trivial identities are applied. If all particles are found somewhere in a measurement region which is beyond the region where the computation takes place, then the quantum computation has been completed. In this model the measurement region is again a fully-2D region.

This disadvantage of not knowing where the computation is at a point in time thus implies some qubit overhead, but perhaps worse, it requires being able to read-out 2D regions of qubits simultaneously. Such read-out is a standard requirement in the usual stationary-qubit circuit model and it necessitates using the third dimension to accomodate the hardware of read-out resonators, measurement feedlines [48], which adds unwanted complexity.

A modification of the rotated lattice geometry of [10] would be to terminate the lattice region in which the computation is followed by an MERA-like expanding *trap* region of $O(1)$ or at most $r = O(\log N)$ sites wide, see Fig. 23. In the trap region new dummy pairs of transmon qubits are placed with their excitations, as shown in Fig. 23. These new dummy excitations can only move forward deeper into the trap region and spawn the forward motion of a next layer of dummy qubits. In the trap region the string is rapidly increasing in length and one expects that this expansion will lead to an irreversible effect on the dynamics, i.e. entropically trap the string, while the overall dynamics remains unitary. It is an interesting open question how to analyze the dynamics of the space-time circuit-to-Hamiltonian construction applied to a MERA circuit (see [49]) or a regular circuit appended by a MERA-like trap region.

One imagines that in such scheme, only the transmon qubits in the trap region whose causal history involves the computational region C need to be measured ¶. Since the trap region is, say, $O(1)$ in width, it would imply that space may be available to let all read-out lines come in from the right side of the planar lattice.

For the Feynman single-clock model, we discuss a different way of making sure that the computation arrives in a certain time due to Peres [50] (and used in [51] in the context of perfect quantum state transfer) in Appendix B.1. Since this solution is not scalable to large system sizes and we don't know how to apply it to the multi-clock model, we do not advocate it as the preferred solution.

¶ One could even imagine applying CNOTs between the computational qubits and the dummy qubits so that the output of the computation is redundantly copied onto the dummy qubits and all qubits are measured.

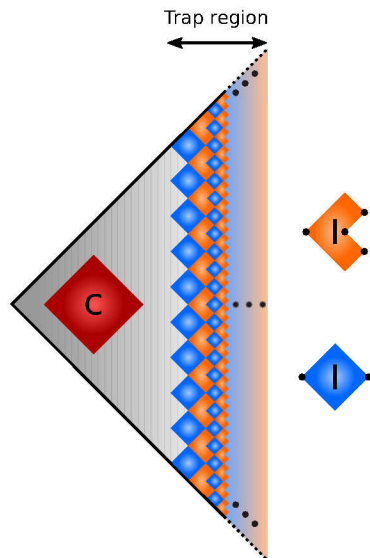


Figure 23: Sketch of a lattice terminated with an entropic trap region. The red square region is where the computation takes place, meaning that non-trivial gates are applied. All blue plaquettes beyond this region execute I gates. In the trap region, new tracks with particles get added in each layer, e.g. an orange ‘isometry’ unit maps a single input track with one particle onto 3 tracks each with a particle. The isometry could otherwise execute the I gate on the input particle. Translated to transmon qubits, we add two pairs of transmon qubits, each with a single excitation, for an orange plaquette. These excitations can only move forward in the right direction and spawn new excitations until they hit the trap end.

6. Discussion

We have observed that adding disorder in the on-site or hopping terms can lead to a localization of time, i.e. the average position of the particle on the middle track can become frozen. It would be interesting to obtain more theoretical or experimental evidence for this behavior for larger lattice sizes N . A difficulty is that disorder breaks the elegant map from the 2D rotated lattice model onto that of $N - 1$ non-interacting fermions hopping on a line, see e.g. [10]. In the non-interacting fermion model the on-site energy of one fermion will now depend on how many fermions are to the left of the site where the fermion sits. Hence a new theoretical analysis, going beyond an Anderson localization analysis [52], may be required to address this question.

We did not explicitly examine the effect of disorder in the execution of logical gates such as the Hadamard or X gate. Since the hopping strengths J determine the unitary gate which is executed, it is clear that such disorder leads to inaccurate computation. Fascinating is that static disorder alters the application of unitary gates U and $U^\dagger$ to possibly irreversible gates M and $M^\dagger$ which could lead to inadvertently measuring the computational state.

An interesting open question is whether the wavefront dynamics can be viewed as

a non-classical ‘bunching’ effect of photons in the following sense. Instead of a pair of transmon qubits we could use two bosonic modes and encode $|0\rangle \equiv |\alpha\rangle_0 |\text{vac}\rangle_1$ and $|1\rangle \equiv |\text{vac}\rangle_0 |\alpha\rangle_1$ where $|\alpha\rangle$ is a coherent state. We can view all flip-flop interactions in our scheme as bosonic beam-splitter interactions, evolving the coherent amplitudes at two adjacent sites on a single track into new coherent amplitudes. Since we have non-linear bosonic interactions, we cannot efficiently solve the system’s dynamics, although a mean-field approximation of the non-linearity would enable an efficient simulation. It will be interesting to understand how this is different from our single-excitation encoding, since the dual-rail coherent state encoding would suffer less from excitation loss. We expect that the coherent state encoding would not lead to wavefront dynamics, nor to the correct application of logical gates.

One could envision incorporating quantum error correction, include the initialization of ancilla qubits during the computation, but this may require going to a 3D version of the model. A mathematical version of such 3D model is discussed in [22]: in this model the computation would move forward as a surface of a crystal which is growing from a corner. Going to 3D means a higher-connectivity per (transmon) degree of freedom and facing practical challenges related to 3D superconducting qubit hardware [48]. In addition, error correction requires the inclusion of ancilla qubits during the computation which are measured and reset to remove entropy build-up, hindering the overall passivity of the scheme.

BMT acknowledges support from the European Research Council (EQEC, ERC Consolidator Grant No. 682726). DD was supported by Intelligence Advanced Research Projects Activity (IARPA) under contract W911NF-16-0114. AC acknowledges support from the Excellence Initiative of the Deutsche Forschungsgemeinschaft. We thank Fabian Hassler, Martin Rymarz and Stefano Bosco for useful discussions.

A. Hamiltonian for the direct Toffoli Gate

We describe the ideas behind the direct Toffoli gate mathematically. We consider how the Hamiltonian H_{valid} is modified as compared to the standard case, Eq. (2). We need to modify the edge Hamiltonian corresponding to the green edges in Fig. 10. We call the set of these edges E_{Tof} . H_{valid} can be written as

$$H_{\text{valid}} = -\Delta \sum_{(\mu,\nu) \in E \setminus E_{\text{Tof}}} \mathbf{n}[\mu] \mathbf{n}[\nu] - \Delta \sum_{e \in E_{\text{Tof}}} H_e.$$

The terms for the edges c_1, d_1 and a_2, b_2 do not change as compared to the standard case, hence we write

$$H_{c_1} = \mathbf{n}[i-1, j] \mathbf{n}[i, j], \quad (\text{A.1a})$$

$$H_{d_1} = \mathbf{n}[i, j+1] \mathbf{n}[i, j], \quad (\text{A.1b})$$

$$H_{a_2} = \mathbf{n}[i+1, j] \mathbf{n}[i+1, j+1], \quad (\text{A.1c})$$

$$H_{b_2} = \mathbf{n}[i+2, j+1] \mathbf{n}[i+1, j+1]. \quad (\text{A.1d})$$

The remaining edges in E_{Tof} are chosen as

$$H_{a_1} = n_{s_1=0}[i, j-1] \mathbf{n}[i, j, 0] + n_{s_1=1}[i, j-1] \mathbf{n}[i, j, 1], \quad (\text{A.2a})$$

$$H_{b_1} = n_{s_1=0}[i+1, j] \mathbf{n}[i, j, 0] + n_{s_1=1}[i+1, j] \mathbf{n}[i, j, 1], \quad (\text{A.2b})$$

$$H_{c_2} = n_{s_2=0}[i, j+1] \mathbf{n}[i+1, j+1, 0] + n_{s_2=1}[i, j+1] \mathbf{n}[i+1, j+1, 1], \quad (\text{A.2c})$$

and

$$H_{d_2} = n_{s_2=0}[i+1, j+2] \mathbf{n}[i+1, j+1, 0] + n_{s_2=1}[i+1, j+2] \mathbf{n}[i+1, j+1, 1]. \quad (\text{A.2d})$$

As for the CNOT one can check that valid strings have energy $E_0 = -(N_{\text{track}} - 1)\Delta$, while the gap still remains Δ . In particular, the first excited subspace of the new H_{valid} is formed by the subspace of strings that are broken at one position and strings that are connected but incorrect.

The hopping Hamiltonian V_{hop} needs to be modified at plaquettes $p_1, p_2, p_3 \in P_{\text{Tof}}$ in Fig. 10. V_{hop} is written as

$$V_{\text{hop}} = -J \sum_{p \in P \setminus P_{\text{Tof}}} V_{\text{hop}, p} - J \sum_{p \in P_{\text{Tof}}} \tilde{V}_{\text{hop}, p},$$

where the three terms in the last sum are given by

$$\tilde{V}_{\text{hop}, p_1} = \sum_{s_1=0}^1 \sum_{k=0}^1 (a_{s_1}^\dagger[i, j, k] a_{s_1}[i-1, j-1] + \text{h.c.}), \quad (\text{A.3a})$$

$$\begin{aligned} \tilde{V}_{\text{hop}, p_2} = & \sum_{s_2=0}^1 (a_{s_2}^\dagger[i+1, j+1, 0] a_{s_2}[i, j, 0] + \text{h.c.}) + \sum_{s_2=0}^1 (a_{s_2}^\dagger[i+1, j+1, 0] a_{s_2}[i, j, 1] + \text{h.c.}) + \\ & \sum_{s_2=0}^1 (a_{s_2}^\dagger[i+1, j+1, 1] a_{s_2}[i, j, 0] + \text{h.c.}) + \sum_{s_2=0}^1 (a_{s_2}^\dagger[i+1, j+1, 1] a_{s_2}[i, j, 1] + \text{h.c.}), \end{aligned} \quad (\text{A.3b})$$

and

$$\tilde{V}_{\text{hop}, p_3} = \sum_{s_2=0}^1 \sum_{k=0}^1 (a_{s_2}^\dagger[i+2, j+2] a_{s_2}[i+1, j+1, k] + \text{h.c.}). \quad (\text{A.3c})$$

To obtain the effective Hamiltonian, we can invoke Schrieffer-Wolff perturbation theory (as in [11]), although in our lowest-order application its effect is rather immediate. For SW theory, we can use Sec. 3 of [53] and Sec. 4 of [54]. Let us denote by P_- the projector onto the groundspace of H_{valid} , i.e., the subspace of valid strings. We also define $P_+ = I - P_-$ the projector onto the subspace of invalid (either disconnected and/or incorrect) strings. Given a generic linear operator O we define the operators

$$\begin{aligned} O_{--} &= P_- O P_- \quad , \quad O_{++} = P_+ O P_+ \\ O_{+-} &= P_+ O P_- \quad , \quad O_{-+} = P_- O P_+. \end{aligned}$$

An operator O is said to be block-diagonal if $O_{+-} = O_{-+} = 0$, while block-off-diagonal if $O_{++} = O_{--} = 0$. Assuming $2\|V\|/\Delta < 1$, the effective low energy Hamiltonian is defined in general as

$$H_{\text{eff}} = (e^S H e^{-S})_{--},$$

with S an anti-hermitian and block-off-diagonal operator such that $e^S H e^{-S}$ is also block-diagonal and $\|S\| < \pi/2$. The operator S can be Taylor-expanded and a recursive relation can be obtained for each power, but here we only use the lowest-order term $S \approx 0$ so the SW transformation $e^S \approx I$. In lowest-order one simply has

$$H_{\text{eff}} = H_{--} + \mathcal{O}(\|V_{\text{hop}}\|^2/\Delta) = (V_{\text{hop}})_{--} + \mathcal{O}(\|V_{\text{hop}}\|^2/\Delta).$$

For the Toffoli region, we can write the effective Hamiltonian as

$$H_{\text{eff}} = -J \sum_{p \in P \setminus P_{\text{Tof}}} H_{\text{cond.hop},p} - J \sum_{p \in P_{\text{Tof}}} \tilde{H}_{\text{cond.hop},p} + \mathcal{O}(\|V_{\text{hop}}\|^2/\Delta),$$

where the second sum involves the conditional hopping terms related to plaquettes $p_1, p_2, p_3 \in P_{\text{Tof}}$. We have

$$\tilde{H}_{\text{cond.hop},p_1} = \sum_{s_1=0}^1 \left(\sum_{s=0}^1 n_{s_1}[i, j-1] \mathbf{n}[i-1, j] a_s^\dagger[i, j, s_1] a_s[i-1, j-1] + \text{h.c.} \right)$$

The term corresponding to the second plaquette is more complicated: on the valid subspace, it acts as

$$\tilde{H}_{\text{cond.hop},p_2} = \sum_{s_1=0}^1 \sum_{s_2=0}^1 \sum_{s=0}^1 n_{s_2}[i, j+1] n_{s_1}[i+1, j] \left(a_{s \oplus s_1 s_2}^\dagger[i+1, j+1, s_2] a_s[i, j, s_1] + \text{h.c.} \right),$$

exhibiting the Toffoli gate logic on the internal spin s . Analogously to the term for the plaquette p_1 , we have

$$\tilde{H}_{\text{cond.hop},p_3} = \sum_{s_2=0,1} \sum_{s=0}^1 n_{s_2}[i+1, j+2] \mathbf{n}[i+2, j+1] \left(a_s^\dagger[i+2, j+2] a_s[i+1, j+1, s_2] + \text{h.c.} \right).$$

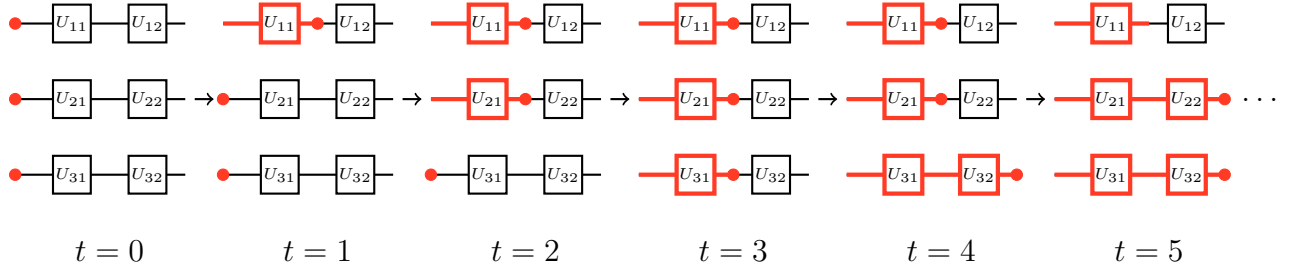


Figure B1: Example of depth-2 quantum circuit with 3 qubits in which gates are executed in a ‘snake-like’ order.

To probe our understanding of the Toffoli construction, we can ask if all the hopping terms described above are really necessary. Let us consider again Fig. 10 and suppose that the first control qubit is in the $s_1 = 1$ state. The target qubit will then be directed to the site $(i, j, 0)$. At this point we are tempted to say that no matter what the state of the second control is, the identity has to be applied and so the hopping from $(i, j, 0)$ to $(i + 1, j + 1, 1)$ looks useless. This is however not the case, and the reason why this is incorrect is that we selected a particular direction of the computation, i.e., from left to right, while terms that hop in the opposite direction are always present by the Hermiticity of the Hamiltonian.

If we remove the identity connecting $(i, j, 0)$ and $(i + 1, j + 1, 1)$ and try to run the circuit backwards one can show that it cannot implement a Toffoli gate. Suppose that the target particle is in $(i + 2, j + 2)$ and the second control is in the state $s_2 = 1$. If the target particle runs backwards, it would be directed to site $(i + 1, j + 1, 1)$. However, at this point even if the first control qubit is in $s_1 = 0$, the target can never go further backwards with the application of the identity since the hopping from $(i + 1, j + 1, 1)$ to $(i, j, 0)$ is missing. We conclude that if the particles hop backwards the (inverse) Toffoli gate is not applied. For this reason, the interactions depicted in Fig. 10 and written in Eqs. (A.1), (A.2) and (A.3) are all necessary. As a final remark, and in analogy with the CNOT, the control particles 1 and 2 must obviously undergo the identity in the Toffoli region.

B. Feynman model with hopping particles on a 2D lattice

We construct a lattice model similar to that in Fig. 6, but now the dynamics of the valid strings maps to that of a quantum walk on a line. This is the dynamics of the clock Hamiltonian proposed by Feynman [1]. A known difference between the multi-clock (as exemplified in the Lloyd-Terhal model) versus this single-clock Feynman construction is that the former allows for a partially parallel, and hence faster, execution of gates.

Before introducing the new model, we describe the type of quantum circuit that we are going to implement. As in Sec. 2 we first focus for simplicity on the case in which only single-qubit gates are present and then extend the analysis in Section B.2 to circuits in which controlled two- and three-qubit gates need to be executed.

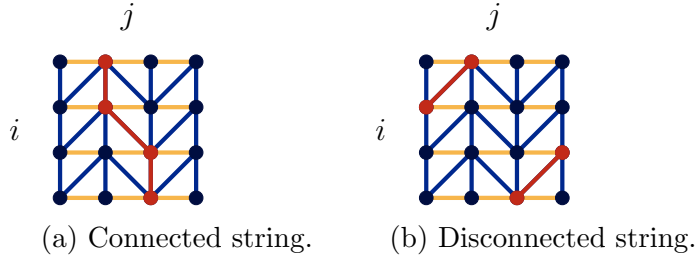


Figure B2: Examples of connected and disconnected strings of particles in the lattice model with $N \times M$ sites (shown is $N = M = 4$). The red dots denote the position of the particles at sites (i, j) .

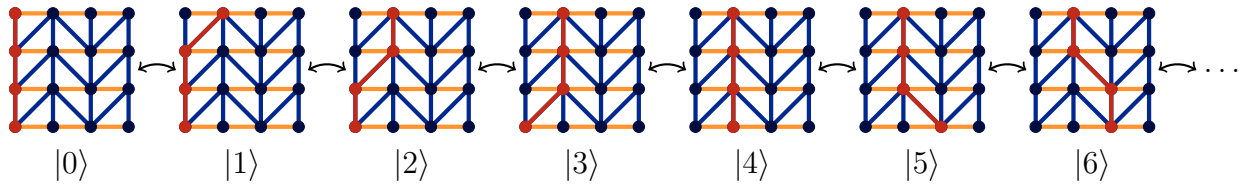


Figure B3: Motion of the string of particles which keeps the string connected and executes the gates in the circuit in the correct order.

Consider a quantum circuit with single-qubit gates shown in Fig. B1. In practice, odd columns in the circuit are executed from top to bottom, while even columns are executed from bottom to top. The motion of the cursor reminds one of a snake that runs over the quantum circuit: this construction has previously been used in e.g. [55, 56]. The reason why we restrict the analysis to this kind of ordering in the execution of the gates is to keep the interactions local on a 2D lattice.

We consider the lattice in Fig. B2, consisting of a grid of sites (i, j) with $i = 1, \dots, N, j = 1, \dots, M$, where particles can reside. Like for the model described in Sec. 2, we assume that there is one particle per horizontal line or track. The model encodes a quantum circuit with N qubits, where each qubit undergoes $M - 1$ unitary gates. The blue edges on the lattice represent *attractive* interactions whose purpose is to keep the string of particles together. Note that, as in the Lloyd-Terhal model, each site participates in 4 attractive edges. The orange edges are associated with hopping terms which apply gates to the qubit that lives on each track as in Eq. (4) in the main text. We denote the set of blue edges as E_b and the set of orange edges as E_y . A generic unitary of our quantum circuit applied at an orange edge e is denoted as U_{ij} as shown in Fig. B1.

Annihilation, creation and number operators are used in the same way as in Sec. 2. We define the Hamiltonian $H = H_{\text{valid}} + V_{\text{hop}}$ as in Eq. (1) with

$$H_{\text{valid}} = -\Delta \sum_{(\mu, \nu) \in E_b} \mathbf{n}[\mu] \mathbf{n}[\nu].$$

The groundspace of this Hamiltonian has eigenvalue $E_0 = -\Delta(N - 1)$ and consists of

all possible strings of particles connected by the blue edges. The number of connected strings in this subspace is $L \equiv N \times (M - 1) + 1$. The first excited space is the space of strings that fail to be connected by blue edges at one position, having energy $E_0 + \Delta$.

The hopping Hamiltonian equals

$$V_{\text{hop}} = -J \sum_{e \in E_y} V_{\text{hop},e},$$

with the term for edge $e = (\mu, \nu)$ between sites $\mu = (i, j + 1)$ and $\nu = (i, j)$ equal to

$$V_{\text{hop},e} = \sum_{s=0}^1 \sum_{s'=0}^1 \langle s' | U_{ij} | s \rangle a_{s'}^\dagger[i, j + 1] a_s[i, j] + \text{h.c.}$$

The effective conditional-hopping Hamiltonian H_{eff} induces a quantum walk on a line with L sites, where the quantum computation follows the order depicted in Fig. B1. An example of this effective motion is depicted in Fig. B3. The system is initialized in a configuration in which all particles are on the left, denoted as $|0\rangle$. From this configuration there is only one *move* that keeps the string connected which brings the string to configuration $|1\rangle$. From here there are two moves that keep the string connected. Either the string hops back to $|0\rangle$ or it hops forward to $|2\rangle$. The conditional hopping proceeds in this way until we reach the bottom, configuration $|4\rangle$, and then we can move backward or forward until we reach the top again etc. We understand that this is exactly the order in which the gates should be implemented

Mathematically, the effective Hamiltonian H_{eff} is given by the projection of H onto the groundspace of connected strings. Each connected string can be labeled as $|k\rangle \in \mathcal{H}_C$ with $k \in \{0, 1, 2, \dots, L-1\}$ as in Fig. B3 in tensor product with the internal state space of the particles. We can thus write

$$H_{\text{eff}} = H_F = -J \sum_{k=0}^{L-1} U_k \otimes |k+1\rangle \langle k| + \text{h.c.},$$

where U_k is the k -th single-qubit gate in our snake-like ordering. Properties of this effective Feynman Hamiltonian H_F are well-known: it performs a unitary evolution in a computational space spanned by the orthogonal vectors

$$|\Psi_k\rangle = (U_k U_{k-1} \dots U_1 |\psi_0\rangle) \otimes |k\rangle = |\psi_k\rangle \otimes |k\rangle,$$

where $|\psi_0\rangle$ is some arbitrary initial state. In the clock subspace $\mathcal{H}_C$ it acts as

$$H_F|_{\mathcal{H}_C} = -J \sum_{k=0}^{L-1} |\Psi_{k+1}\rangle \langle \Psi_k| + \text{h.c.},$$

which up to a trivial relabelling $|\Psi_k\rangle \equiv |k+1\rangle$ is just the Hamiltonian of a continuous-time quantum walk on a line with L sites, i.e.

$$H_F|_{\mathcal{H}_C} = H_L = -J \sum_{k=1}^L |k+1\rangle \langle k| + \text{h.c.} \quad (\text{B.1})$$

Diagonalizing this Hamiltonian H_L given eigenstates [51, 26]

$$|\tilde{l}\rangle = \sqrt{\frac{2}{L+1}} \sum_{k=1}^L \sin\left(\frac{\pi lk}{L+1}\right) |k\rangle$$

and corresponding eigenvalues

$$\tilde{\omega}_l = -2J \cos\left(\frac{\pi l}{L+1}\right),$$

for $l = 0, \dots, L-1$.

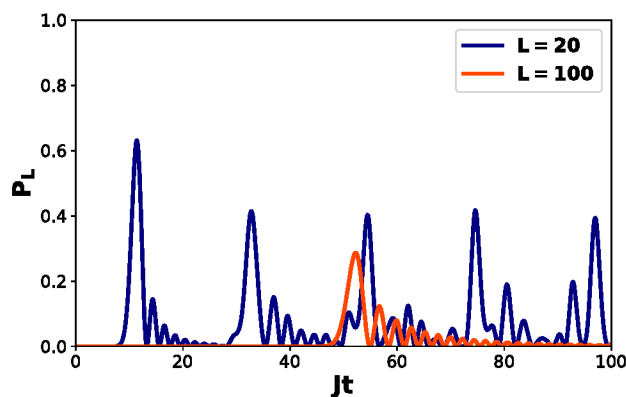


Figure B4: Success probability $P_L(t)$ for $L = 20$ and $L = 100$ as a function of Jt .

The probability that starting from the initial computational state $|1\rangle$ the system is found in the final computational state $|L\rangle$ at some fixed time t equals

$$P_L(t) = |\langle L | e^{-iH_L t} | 1 \rangle|^2.$$

with

$$\langle L | e^{-iH_L t} | 1 \rangle = \sum_{l=1}^L \langle L | \tilde{l} \rangle \langle \tilde{l} | e^{-iH_L t} | 1 \rangle = \frac{2}{L+1} \sum_{l=1}^L \sin\left(\frac{\pi l}{L+1}\right) \sin\left(\frac{\pi Ll}{L+1}\right) e^{-i\tilde{\omega}_l t}.$$

As a concrete example, we can plot in Fig. B4 the success probability for the case $L = 20$ and $L = 100$. As expected, we see how the probability decreases as a function of L . For instance, the maximum probability with $L = 1000$ is approximately 0.1.

B.1. Peres' trick

One can use a trick to reach the final stage of computation exactly. In 1985 Peres formulated a modification of the Feynman Hamiltonian which ensures that the computation is executed in a specific time [50]. This trick was rediscovered in [51], in the context of perfect quantum state transfer. It has also been argued that this way of achieving perfect state transfer is optimal from several points of view [57].

The idea is to choose different hopping strengths J_k so that the Hamiltonian is proportional to the angular momentum operator $\mathcal{J}_x$ of a fictitious spin- $(L-1)/2$ particle. The basis states on the line $|k\rangle$ will be interpreted as the eigenkets of the operator $\mathcal{J}^2$ and $\mathcal{J}_z$. To this end, we relabel the basis states (again) as

$$|k\rangle \rightarrow |m = -(L-1)/2 + k - 1\rangle, \quad k \in \{1, \dots, L\}.$$

which are eigenstates of $\mathcal{J}_z$ with eigenvalue m ranging from $-\frac{L-1}{2}$ to $\frac{L-1}{2}$. The Peres-Feynman Hamiltonian in this basis is then chosen as

$$H_{L,PF} = - \sum_{m=-(L-1)/2}^{(L-1)/2-1} J_m |m+1\rangle \langle m| + \text{h.c.} \equiv -J\mathcal{J}_x$$

with the choice

$$J_m = \frac{J}{2} \sqrt{m(L-m)}$$

Using this Hamiltonian it can be shown that at time $t = T = \pi/J$ we have the desired condition $|\psi(T)\rangle = |L\rangle$ with probability 1.

Note that the coupling parameters J_m scale with the number of gates. For this reason we cannot envision using this trick for a quantum computation of very long length.

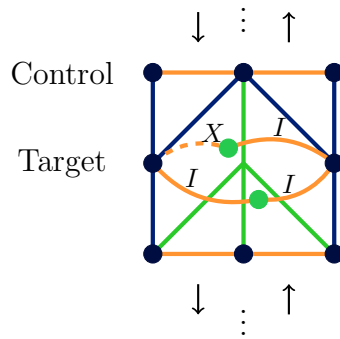
B.2. Multi-Qubit Logic

The construction of the CNOT and Toffoli gates is similar to the construction discussed in Sec. 2.1. We will show that there are two ways of approaching the problem to realize a CNOT gate depicted in Fig. B5 .

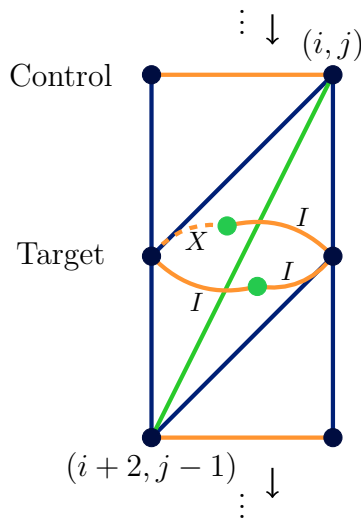
In Fig. B5a we represent a slow CNOT. In this case we introduce a split-site in the lattice and we need to modify the green edges in the Hamiltonian (as compared to the standard case) and require the hopping terms to the split-sites to enact a possible X -gate. This is essentially the same gadget as the CNOT gadget in the Lloyd-Terhal scheme, described in Section 2.1.1. In this construction, the CNOT *begins* in one column, but it is only terminated in the next one. In the meantime all the gates below the target in the first column and in the second column need to be executed before the CNOT can be finally terminated. This means that the control and the target particles are *busy* for two rounds of the computation and cannot be involved in other gates.

To avoid this problem, we might consider implementing the CNOT as in Fig. B5b, which we call a fast CNOT since it can be accomplished in a single round. The idea is to add a new split-site for the CNOT in the middle of an orange hopping edge. The horizontal site coordinate j can be given half-integer values at such a split-site, i.e. one uses annihilation operators $a_s[i, j + 1/2, k]$ with $k = 0, 1$ labeling the two sites. We add new attractive green edges to this split-site which selectively depend on the state of the control qubit, e.g. the green edges in Fig. B5b correspond to

$$-\Delta \sum_{s=0,1} n_s[i, j] \mathbf{n}[i+1, j-1/2, s] - \Delta \mathbf{n}[i+2, j-1] \mathbf{n}[i+1, j-1/2].$$



(a) Slow CNOT.



(b) Fast CNOT

Figure B5: CNOT in the ‘single-clock snake’ lattice model of Fig. B2. The green edges are those who need to be modified similar to what we do in Subsec. 2.1.1.

We can check that the CNOT is fully executed before the computation continues to run downwards as in Fig. B5b. Note that the control particle can be either be above or below the target particle. This fast construction is advantageous since each CNOT will just occupy one site in the quantum walk on the line, while this conclusion cannot be drawn for the slow CNOT, as its effect in terms of overhead is algorithm-dependent. The disadvantage of the fast CNOT is that it requires a region which is more dense with interactions. An analogous construction can be imagined for the Toffoli gate. In the fast version each Toffoli adds two sites to the quantum walk. Even though the green split-sites only participate in two attractive edges each, the fast CNOT requires two normal sites which participate in 5 attractive edges in total which adds complexity.

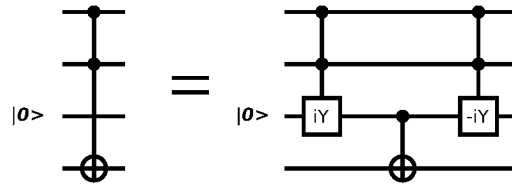


Figure B6: Implementation of a Toffoli gate using an ancilla qubit, a CNOT and controlled-controlled- iY and controlled-controlled- $(iY)^\dagger$ gates.

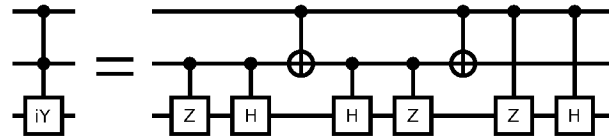


Figure B7: Implementation of a controlled-controlled- iY gate using CNOT, controlled-Hadamard and controlled-Z gates. The implementation of controlled-controlled- $(iY)^\dagger$ follows analogously.

C. Universality of Hadamard, CNOT and controlled-Hadamard

In this Appendix we show how the Toffoli gate can be constructed using Hadamard, CNOT and controlled-Hadamard gates with the help of an ancilla qubit initialized in $|0\rangle$. The key to this construction is shown in Fig. B6, where the Toffoli gate is obtained using a CNOT gate, a controlled-controlled- iY and its Hermitian conjugate.

Now we show that the controlled-controlled- iY , and its Hermitian conjugate can be obtained from Hadamard, CNOT and controlled-Hadamard. To this end we follow the general construction of a controlled-controlled- U described in Sec. 4.3 of Ref. [23]. In order to use this method to construct the controlled-controlled- iY operator, we have to find a unitary operator V such that $V^2 = iY$, which turns out to be the real-valued gate $V = ZH$. As shown in Fig. B7 the controlled-controlled- iY is thus obtained in terms of

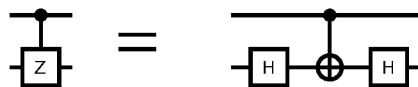


Figure C1: Controlled-Z gate from Hadamard and CNOT.

CNOT, controlled-Hadamard and controlled-Z gate (CZ). Finally, using the well-known relation between the CZ gate and the CNOT gate shown in Fig. C1, we complete the synthesis of a Toffoli gate using Hadamard, CNOT and controlled-Hadamard, showing the universality of this set of quantum gates.

- [1] R. P. Feynman, “Quantum mechanical computers,” *Foundations of Physics*, vol. 16, p. 507, 1986.
- [2] A. Y. Kitaev, A. H. Shen, and M. N. Vyalı, *Classical and Quantum Computation*. Boston, MA, USA: American Mathematical Society, 2002.
- [3] D. Aharonov, W. van Dam, J. Kempe, Z. Landau, S. Lloyd, and O. Regev, “Adiabatic quantum computation is equivalent to standard quantum computation,” *SIAM J. Comput.*, vol. 37, pp. 166–194, Apr. 2007.
- [4] A. Mizel, D. A. Lidar, and M. Mitchell, “Simple proof of equivalence between adiabatic quantum computation and the circuit model,” *Phys. Rev. Lett.*, vol. 99, p. 070502, Aug 2007.
- [5] N. Margolus, “Parallel quantum computation,” in *Complexity, Entropy, and the Physics of Information, SFI Studies in the Sciences of Complexity*, pp. 273–287, Addison-Wesley, 1990.
- [6] D. Janzing, “Spin-1/2 particles moving on a two-dimensional lattice with nearest-neighbor interactions can realize an autonomous quantum computer,” *Phys. Rev. A*, vol. 75, p. 012307, Jan 2007.
- [7] A. Mizel, M. W. Mitchell, and M. L. Cohen, “Energy barrier to decoherence,” *Phys. Rev. A*, vol. 63, p. 040302, Mar 2001.
- [8] A. Mizel, “Mimicking time evolution within a quantum ground state: Ground-state quantum computation, cloning, and teleportation,” *Phys. Rev. A*, vol. 70, p. 012304, Jul 2004.
- [9] N. P. Breuckmann and B. M. Terhal, “Space-time circuit-to-Hamiltonian construction and its applications,” *Journal of Physics A: Mathematical and Theoretical*, vol. 47, no. 19, p. 195304, 2014.
- [10] D. Gosset, B. M. Terhal, and A. Vershynina, “Universal adiabatic quantum computation via the space-time circuit-to-Hamiltonian construction,” *Phys. Rev. Lett.*, vol. 114, p. 140501, Apr 2015.
- [11] S. Lloyd and B. M. Terhal, “Adiabatic and Hamiltonian computing on a 2D lattice with simple two-qubit interactions,” *New Journal of Physics*, vol. 18, no. 2, p. 023042, 2016.
- [12] A. M. Childs, D. Gosset, and Z. Webb, “Universal computation by multiparticle quantum walk,” *Science*, vol. 339, no. 6121, pp. 791–794, 2013.
- [13] A. Childs, D. Gosset, D. Nagaj, M. Raha, and Z. Webb, “Momentum switches,” *Quant. Inf. Comp.*, vol. 15, pp. 0601–0621, 2015.
- [14] Y. Lahini, G. R. Steinbrecher, A. D. Bookatz, and D. Englund, “Quantum logic using correlated one-dimensional quantum walks,” *npj Quantum Information*, vol. 4, no. 1, p. 2, 2018.
- [15] I. Georgescu, S. Ashhab, and F. Nori, “Quantum simulation,” *Rev. Mod. Phys.*, vol. 86, pp. 152–185, Mar 2013.
- [16] J. Koch, T. M. Yu, J. Gambetta, A. A. Houck, D. I. Schuster, J. Majer, A. Blais, M. H. Devoret, S. M. Girvin, and R. J. Schoelkopf, “Charge-insensitive qubit design derived from the Cooper pair box,” *Phys. Rev. A*, vol. 76, p. 042319, Oct 2007.
- [17] S. Richer, N. Maleeva, S. T. Skacel, I. M. Pop, and D. DiVincenzo, “Inductively shunted transmon qubit with tunable transverse and longitudinal coupling,” *Phys. Rev. B*, vol. 96, p. 174520, Nov 2017.
- [18] N. E. Frattini, U. Vool, S. Shankar, A. Narla, K. M. Sliwa, and M. H. Devoret, “3-wave mixing Josephson dipole element,” *Applied Physics Letters*, vol. 110, no. 22, p. 222603, 2017.
- [19] M. Kounalakis, C. Dickel, A. Bruno, N. K. Langford, and G. A. Steele, “Tuneable hopping and nonlinear cross-Kerr interactions in a high-coherence superconducting circuit,” *npj Quantum Information*, vol. 4, 2018.
- [20] M. C. Collodo, A. Potočník, S. Gasparinetti, J.-C. Besse, M. Pechal, M. Sameti, M. J. Hartmann, A. Wallraff, and C. Eichler, “Observation of the Crossover from Photon Ordering to Delocalization in Tunably Coupled Resonators,” *ArXiv e-prints*, Aug. 2018.
- [21] S. E. Nigg, H. Paik, B. Vlastakis, G. Kirchmair, S. Shankar, L. Frunzio, M. H. Devoret, R. J. Schoelkopf, and S. M. Girvin, “Black-Box Superconducting Circuit Quantization,” *Physical Review Letters*, vol. 108, p. 240502, June 2012.
- [22] R. Dijkgraaf, D. Orlando, and S. Reffert, “Quantum crystals and spin chains,” *Nuclear Physics B*, vol. 811, pp. 463–490, Apr. 2009.

- [23] M. A. Nielsen and I. L. Chuang, *Quantum Information and Quantum Computation*. Cambridge, England: Cambridge University Press, 2000.
- [24] Y. Shi, “Both Toffoli and controlled-NOT need little help to do universal quantum computing,” *Quantum Info. Comput.*, vol. 3, pp. 84–92, Jan. 2003.
- [25] L. Bianchi, N. Pancotti, and S. Bose, “Quantum gate learning in qubit networks: Toffoli gate without time-dependent control,” *Npj Quantum Information*, vol. 2, p. 16019, 2016.
- [26] D. Nagaj, “Fast universal quantum computation with railroad-switch local Hamiltonians,” *Journal of Mathematical Physics*, vol. 51, no. 6, p. 062201, 2010.
- [27] O. V. Marchukov, A. G. Volosniev, D. Valiente, M. Petrosyan, and N. T. Zinner, “Quantum spin transistor with a Heisenberg spin chain,” *Nature Communications*, vol. 7, p. 13070, 2016.
- [28] N. J. S. Loft, L. B. Kristensen, C. K. Andersen, and N. T. Zinner, “Quantum spin transistors in superconducting circuits,” *ArXiv e-prints*, Feb. 2018.
- [29] J. Braumüller, M. Sandberg, M. R. Vissers, A. Schneider, S. Schlögl, L. Grnhaupt, H. Rotzinger, M. Marthaler, A. Lukashenko, A. Dieter, A. V. Ustinov, M. Weides, and D. P. Pappas, “Concentric transmon qubit featuring fast tunability and an anisotropic magnetic dipole moment,” *Applied Physics Letters*, vol. 108, no. 3, p. 032601, 2016.
- [30] J.-M. Reiner, M. Marthaler, J. Braumüller, M. Weides, and G. Schön, “Emulating the one-dimensional Fermi-Hubbard model by a double chain of qubits,” *Phys. Rev. A*, vol. 94, p. 032338, Sep 2016.
- [31] L. Neumeier, M. Leib, and M. J. Hartmann, “Single-photon transistor in circuit quantum electrodynamics,” *Phys. Rev. Lett.*, vol. 111, p. 063601, Aug 2013.
- [32] J. Jin, D. Rossini, R. Fazio, M. Leib, and M. J. Hartmann, “Photon solid phases in driven arrays of nonlinearly coupled cavities,” *Phys. Rev. Lett.*, vol. 110, p. 163605, Apr 2013.
- [33] D. Marcos, P. Widmer, E. Rico, M. Hafezi, P. Rabl, U.-J. Wiese, and P. Zoller, “Two-dimensional lattice gauge theories with superconducting quantum circuits,” *Annals of Physics*, vol. 351, pp. 634 – 654, 2014.
- [34] M. Leib, P. Zoller, and W. Lechner, “A transmon quantum annealer: decomposing many-body ising constraints into pair interactions,” *Quantum Science and Technology*, vol. 1, no. 1, p. 015008, 2016.
- [35] D. M. Pozar, *Microwave engineering; 3rd ed.* Hoboken, NJ: Wiley, 2005.
- [36] V. E. Manucharyan, *Superinductance*. PhD thesis, Yale University, 2012.
- [37] V. E. Manucharyan, N. A. Masluk, A. Kamal, J. Koch, L. I. Glazman, and M. H. Devoret, “Evidence for coherent quantum phase slips across a Josephson junction array,” *Phys. Rev. B*, vol. 85, p. 024521, Jan 2012.
- [38] G. Rastelli, I. M. Pop, and F. W. J. Hekking, “Quantum phase slips in Josephson junction rings,” *Phys. Rev. B*, vol. 87, p. 174513, May 2013.
- [39] Y. V. Nazarov and Y. M. Blanter, *Quantum Transport: Introduction to Nanoscience*. Cambridge University Press, 2009.
- [40] K. A. Matveev, A. I. Larkin, and L. I. Glazman, “Persistent current in superconducting nanorings,” *Phys. Rev. Lett.*, vol. 89, p. 096802, Aug 2002.
- [41] V. E. Manucharyan, J. Koch, L. I. Glazman, and M. H. Devoret, “Fluxonium: Single Cooper-pair circuit free of charge offsets,” *Science*, vol. 326, no. 5949, pp. 113–116, 2009.
- [42] U. Vool, A. Kou, W. C. Smith, N. E. Frattini, K. Serniak, P. Reinhold, I. M. Pop, S. Shankar, L. Frunzio, S. M. Girvin, and M. H. Devoret, “Driving forbidden transitions in the fluxonium artificial atom,” *Phys. Rev. Applied*, vol. 9, p. 054046, May 2018.
- [43] A. Blais, R.-S. Huang, A. Wallraff, S. M. Girvin, and R. J. Schoelkopf, “Cavity quantum electrodynamics for superconducting electrical circuits: An architecture for quantum computation,” *Phys. Rev. A*, vol. 69, p. 062320, Jun 2004.
- [44] F. Brito, F. Rouxinol, M. D. LaHaye, and A. O. Caldeira, “Testing time reversal symmetry in artificial atoms,” *New Journal of Physics*, vol. 17, no. 7, p. 075002, 2015.
- [45] J. Koch, A. A. Houck, K. L. Hur, and S. M. Girvin, “Time-reversal-symmetry breaking in circuit-

- QED-based photon lattices,” *Phys. Rev. A*, vol. 82, p. 043811, Oct 2010.
- [46] C. Müller, S. Guan, N. Vogt, J. H. Cole, and T. M. Stace, “Passive on-chip superconducting circulator using a ring of tunnel junctions,” *Phys. Rev. Lett.*, vol. 120, p. 213602, May 2018.
 - [47] L. Caha, Z. Landau, and D. Nagaj, “Clocks in Feynman’s computer and Kitaev’s local hamiltonian: Bias, gaps, idling, and pulse tuning,” *Phys. Rev. A*, vol. 97, p. 062306, Jun 2018.
 - [48] J. H. Béjanin, T. G. McConkey, J. R. Rinehart, C. T. Earnest, C. R. H. McRae, D. Shiri, J. D. Bateman, Y. Rohanizadegan, B. Penava, P. Breul, S. Royak, M. Zapatka, A. G. Fowler, and M. Mariantoni, “Three-dimensional wiring for extensible quantum computing: The quantum socket,” *Phys. Rev. Applied*, vol. 6, p. 044010, Oct 2016.
 - [49] F. Metz, “Space-time Circuit-to-Hamiltonian Construction applied to a MERA circuit.” Bachelor Thesis, 2015, RWTH Aachen University, http://www.quantuminfo.physik.rwth-aachen.de/global/show_document.asp?id=aaaaaaaaaanxrlq.
 - [50] A. Peres, “Reversible logic and quantum computers,” *Phys. Rev. A*, vol. 32, pp. 3266–3276, Dec 1985.
 - [51] M. Christandl, N. Datta, A. Ekert, and A. J. Landahl, “Perfect state transfer in quantum spin networks,” *Phys. Rev. Lett.*, vol. 92, p. 187902, May 2004.
 - [52] F. Delyon, H. Kunz, and B. Souillard, “One-dimensional wave equations in disordered media,” *Journal of Physics A: Mathematical and General*, vol. 16, no. 1, p. 25, 1983.
 - [53] S. Bravyi, D. P. DiVincenzo, and D. Loss, “Schrieffer-Wolff transformation for quantum many-body systems,” *Annals of Physics*, vol. 326, no. 10, pp. 2793 – 2826, 2011.
 - [54] S. Bravyi and M. Hastings, “On complexity of the quantum Ising model,” *Communications in Mathematical Physics*, vol. 349, pp. 1–45, Jan 2017.
 - [55] R. Oliveira and B. M. Terhal, “The complexity of quantum spin systems on a two-dimensional square lattice,” *Quantum Info. Comput.*, vol. 8, pp. 900–924, Nov. 2008.
 - [56] D. Nagaj, “Universal two-body-hamiltonian quantum computing,” *Phys. Rev. A*, vol. 85, p. 032330, Mar 2012.
 - [57] A. Kay, “Perfect, Efficient, State Transfer and its application as a constructive tool,” *International Journal of Quantum Information*, vol. 08, no. 04, pp. 641–676, 2010.