

Topological network-assisted quantum operations in two-dimensions

MSc COSSE Thesis

by

Genya G. Crossman

Student number 5607728

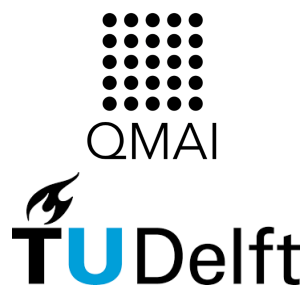
29 August, 2022

Thesis committee

Prof. dr. E. Greplová

Prof. dr. ir. C. Vuik

Prof. dr. dr. ir. AW Heemink



Contents

1	Introduction	1
1.1	Overview	1
1.2	Bra-Ket/Dirac Notation	2
1.3	Bloch Sphere and Single Qubit Gates	3
2	Circuit QED	6
2.1	Overview	6
2.2	Lumped Element Resonators	6
2.3	Transmon and Jaynes-Cummings Hamiltonian	10
3	Topological quantum systems	14
3.1	SSH Model	14
3.2	Time-Reversal Symmetry and Quantum Spin Hall Effect	20
3.3	Quantum Spin Hall Effect in the Classical Realm	21
4	Optimal Control	25
4.1	GRAPE	25
4.2	DRAG	27
5	Implementation	30
5.1	Mean Field Approximation	30
5.2	The Lattice	31
5.3	Pulse Shaping	33
5.4	Metrics	34
6	One Qubit	36
6.1	Edge Modes	36
6.2	Qubit Response	40
6.3	Transmission Coefficient	42
6.4	Qubit-Resonator Coupling	45
7	Two Qubits	46
8	Conclusion	50

1 Introduction

1.1 Overview

Scientific computing and applied mathematics enable the exploration of and, sometimes even, the simplification of complex systems through various optimized modeling and simulation methods. These fields create and utilize computational resources to do so. Many problems, however, are still unsolvable or very difficult to solve with current well-established methods – be it algorithms or computers. Quantum computers were theorized and are now being realized to extend computational abilities. Though commercially available and widely researched, quantum computers today are still more proof-of-concept than “universally” applicable tools. A major part of the problem is these systems are riddled with noise, hence the phrase noisy intermediate scale quantum (NISQ) devices is commonly used to refer to current day devices. There are many approaches to mitigating the effects of noise in quantum systems from both hardware and software perspectives. This work focuses on optimizing the device hardware by combining aspects of two already well-explored systems – superconducting quantum integrated circuits and topological insulators.

The former are relatively easy to fabricate and have become one of the main focuses of today’s industry, including competitive involvement from IBM and Google [1, 2]. The heart of superconducting quantum circuits is found in two circuit elements: the quantum bit (qubit) and the resonator. The qubit manages quantum information processing. The resonator is a simple inductor and capacitor (LC) in parallel and is a well understood classical circuit element. Resonators play a key role in this research. In contrast, topological quantum systems are incredibly difficult to fabricate, having only recently been demonstrated for the first time by Microsoft [3]. In theory, however, topological insulators promise a strong shield against noise across the device [4]. By leveraging the simplicity and ease of simulating classical LC resonators and combining this with the noise-protection introduced in topological systems, this project lays the groundwork for developing a model supporting how such a hybrid hardware could process and protect quantum information [5]. This thesis presents the characterization of how a two-dimensional array of LC resonators could behave similarly to a qubit when demonstrating the topological property known as edge modes (see figure 1) and proposes an approach for how to model such systems.

Understanding how to control these modes, though, requires a familiarity with the theory behind quantum integrated circuits, circuit quantum electrodynamics (CQED) (see chapter 2), as well as the importance of energy band gaps and symmetries in topological systems (see sections 3.1 and 3.2). After an introduction to these topics, section 3.3 presents an overview of experimental evidence of edge modes on pendula (the mechanical equivalent to LC resonators) as motivation for how to demonstrate edge modes in two-dimensional classical resonator systems [6]. Chapter 4 serves as a brief introduction to two common optimal control methods for pulse design: gradient ascent pulse engineering (GRAPE) [7] and derivative removal by adiabatic gate (DRAG) [8]. While

neither were directly applied in this work, the two methods inspired the implemented optimization technique. Chapter 5 introduces the specifics of how the system was simulated. It begins with highlighting the use of the mean field approximation for ease of solving the system numerically and continues with discussing the specifics around the lattice architecture and pulse shaping. Section 5.4 discusses which metrics – namely, the transmission coefficient – should be used to determine the quality of the system. Results on simulations of an array of resonators connected to a single qubit and two qubits are presented and discussed in chapters 6 and 7, respectively. Finally, recommendations for how to continue studying these systems are outlined in chapter 8. Prior to any of the particularities of this work, though, it is useful to become familiar with some standard notation, vocabulary, and visualization methods of quantum computing.

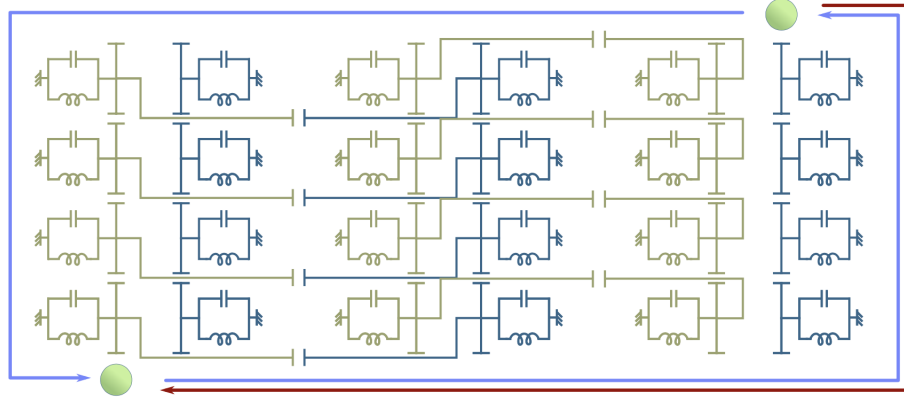


Figure 1: A sketch of the proposed hardware. A two-dimensional array of resonators, which is designed in such a way as to demonstrate topological edge modes, is attached to two control qubits. The two edge modes of the array act as ground and excited states, causing the array to behave like a qubit itself [5]. Image is from Greplová (2021) [5].

1.2 Bra-Ket/Dirac Notation

Quantum mechanics is, in its mathematical essence, an elegant discussion around linear algebra. Quantum states are represented by vectors and operators by matrices. Often times these states and operators are commonly found throughout many different quantum systems, so rather than re-write the lengthy linear algebraic representation, the bra-ket notation is used. This is also known as Dirac notation. In this a vector is written as a ket, shown as

$$\vec{0} = |0\rangle = \begin{bmatrix} 1 \\ 0 \end{bmatrix} \quad (1)$$

This state is known as the ground state of a qubit, whereas

$$|1\rangle = \begin{bmatrix} 0 \\ 1 \end{bmatrix} \quad (2)$$

is the excited state [9]. Note that a qubit only has two energy states of interest – $|0\rangle$ and $|1\rangle$ – in the same way that its classical counterpart, the bit, only considers 0 and 1; if more states are accessible then it is no longer a qubit, e.g. a qutrit actively considers gates on $|0\rangle$, $|1\rangle$, and $|2\rangle$ [9]. When discussing the superconducting qubit called a transmon, as in section 2.3, however, the $|2\rangle$ state is included in some calculations as the energetic boundary whose condition must be considered to understand the two lower energetic states of interest. The conjugate transpose, or Hermitian conjugate, of a ket is known as a bra, or

$$\vec{0}^H = \vec{0}^* = \vec{0}^\dagger = \langle 0| = [1 \quad 0] \quad (3)$$

A bra and ket put immediately next to each other represent an inner product [9]. It is easy to see, then, that $\langle x|x\rangle = 1$ for any state $|x\rangle$ and that $\langle x|y\rangle = 0$ for any orthogonal states $|x\rangle$ and $|y\rangle$. Though it may seem superfluous to recreate a notation for already well-established linear algebraic methods, this notation greatly simplifies equations when considering larger systems beyond a single qubit. For example, given a two qubit system of qubits a and b , the state of the whole system could be written as the tensor product of the two respective states [9].

$$|a\rangle \otimes |b\rangle = |a\rangle |b\rangle = |ab\rangle \quad (4)$$

Operators, of which quantum gates are a subset, are often denoted with a hat and represent matrices. One example of a quantum gate acting on a single qubit is the X gate

$$\hat{X} = \sigma_x = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \quad (5)$$

which will flip the qubit it acts on from $|0\rangle$ to $|1\rangle$ and vice versa. This is also known as the Pauli matrix σ_x and is known in classical computing as the NOT gate [9]. As quantum gates have common labels (e.g. X, Y, Z, H, CNOT, etc.), the hat notation highlighting their responsibility as operators is dropped. The X gate is also defined as a rotation of π about the x-axis of the Bloch sphere, which is useful tool for visualising the evolution of a single qubit's state.

1.3 Bloch Sphere and Single Qubit Gates

If one drew a connection between the two possible values of a classical bit, it would be a line as there are only two possibilities: 0 or 1. Though a qubit can be 0 or 1 (written in Dirac notation that is $|0\rangle$ or $|1\rangle$), it also can be what is called a superposition of the two states, or rather some combination of the two

$$|\psi\rangle = \alpha |0\rangle + \beta |1\rangle \quad (6)$$

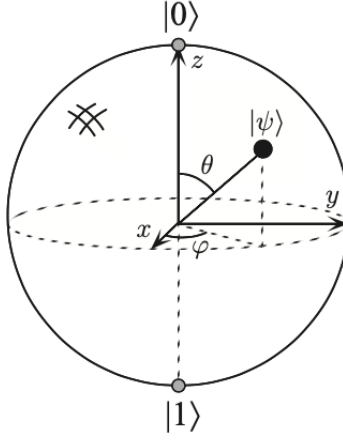


Figure 2: The Bloch sphere helps visualize how the state of a single qubit evolves. The north pole is the ground state and the south pole is the excited state. Where a classical bit could only exist on these poles, the qubit's state, $|\psi\rangle$, can be anywhere on the the surface of the sphere. Image is from Nielson and Chuang (2010) [9].

where $\alpha, \beta \in \mathbb{C}$. The tricky thing with quantum information, however, is that it cannot be directly observed. If one were to measure $|\psi\rangle$ in Equation 6, one would get the result $|0\rangle$ with probability of $|\alpha|^2$ or $|1\rangle$ with probability of $|\beta|^2$ [9]. The probabilities of measurement must sum to 1, $|\alpha|^2 + |\beta|^2 = 1$, implying that $|\psi\rangle$ is normalized, much like a unit vector [9]. When expressing the state in terms of spherical coordinates, any global phase is ignored as it has no physical significance [9].

$$|\psi\rangle = \cos\left(\frac{\theta}{2}\right) |0\rangle + e^{i\varphi} \sin\left(\frac{\theta}{2}\right) |1\rangle \quad (7)$$

θ is the angle from the z-axis and φ is the angle from the x-axis. With this, it is easy to see that qubit state $|\psi\rangle$ can be anywhere on the Bloch sphere's surface (see figure 2). This is a very easy way to see how single qubit gates affect a state. Let's revisit the example of the X gate in equation 5 to see this clearly.

When an X gate is applied on a quantum state, the state flips about the x-axis of the Bloch sphere. For example, if $|\psi\rangle = |0\rangle$ to start and an X gate is applied, then $|\psi\rangle = |1\rangle$.

$$X |0\rangle = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \end{bmatrix} = |1\rangle \quad (8)$$

Similarly, $X |1\rangle = |0\rangle$ [9]. The Z gate is similar and is defined by it's Pauli matrix.

$$Z = \sigma_z = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \quad (9)$$

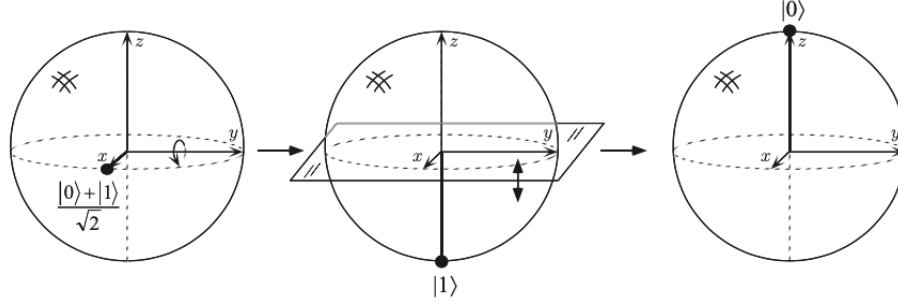


Figure 3: A Hadamard gate acting on initial state $\frac{|0\rangle+|1\rangle}{\sqrt{2}}$ results in final state $|0\rangle$. The evolution of the state is marked by the black dot in each frame. Image is from Nielson and Chuang (2010) [9].

In this case, $Z|0\rangle = 0$ and $Z|1\rangle = -1$ [9]. There is also a Y gate, which behaves in a similar fashion to the X and Z . What is even more interesting, is the Hadamard gate, H , which is known for generating a superposition state.

$$H = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \quad (10)$$

where

$$\begin{aligned} H|0\rangle &= \frac{|0\rangle + |1\rangle}{\sqrt{2}} \\ H|1\rangle &= \frac{|0\rangle - |1\rangle}{\sqrt{2}} \end{aligned} \quad (11)$$

These superposition states are on either end of the x-axis of the Bloch sphere. If a Hadamard gate is run on a qubit that is already in a superposition state directly along the x-axis, for example $\frac{|0\rangle+|1\rangle}{\sqrt{2}}$, then the gate causes the state to return to $|0\rangle$, as is demonstrated in figure 3. In this way, applying two Hadamard gates to initial state $|0\rangle$ results in an identical final state, i.e. $H^2 = \mathbb{I}$ [9]. This is in contrast to the other gates, where $XX^\dagger = ZZ^\dagger = \mathbb{I}$. This touches upon the single rule for quantum gates: they must be unitary [9].

While mapping quantum gates on the Bloch sphere is useful, it does not explain how to actually run gates on quantum hardware. This project studies a model for a new type of hardware in anticipation of implementing single qubit gates – namely the X , Z , and H gates. As the hardware in question is rooted in superconducting qubits, it is useful to understand the foundational theory behind that field, namely, CQED.

2 Circuit QED

2.1 Overview

Superconducting quantum circuits have become one of the central quantum hardware of today largely for the ease in fabrication and scalability. These devices are made based on well-established lithography processes. Nonetheless, it is still difficult to properly model and fabricate these systems, specifically related to a particular circuit element crucial to creating a qubit: the Josephson junction (JJ). JJs determine the frequency of a qubit and, thus, if and how it will interact with its neighbors. This project's array of resonators would bypass this problem. Resonators are easier for classical computers to model and are more robust against fabrication errors compared to JJs. If using a two-dimensional array of resonators as a qubit (rather than the commonly used transmon, which contains JJs) works well, then it would be easier to simulate and build more reliable and larger quantum systems. The motivation for how a block of resonators could demonstrate single qubit gates is rooted in work done by Süssstrunk and Huber (2015) and requires an understanding of topological systems [6]. This will be addressed in the next chapter. Prior to that, though, it is important to understand how resonators themselves work, as well as how a special type of superconducting qubit called a transmon works and how these two elements and their interaction can be described by the Jaynes-Cummings Hamiltonian.

2.2 Lumped Element Resonators

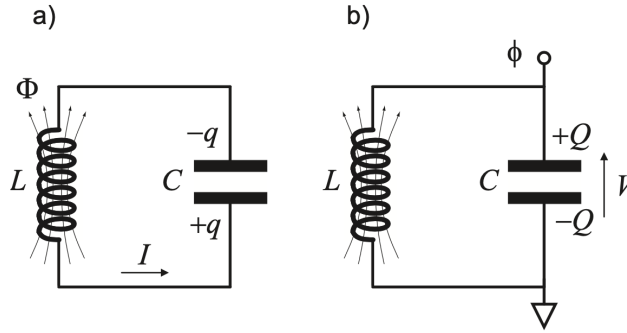


Figure 4: The LC resonator is comprised of an inductor, L , and capacitor, C , in parallel. On the left-hand side, the system is expressed in terms of flux about the inductor, charge across the capacitor, and current, I ; whereas on the right-hand side, the system is explained by a flux, ϕ , at the node north of the capacitor and a voltage, V , across the capacitor. Both diagrams portray the same system, but expressed with different degrees of freedom: charge and flux on the left and right respectively. Image is from Girvin (2014) [10].

Resonators in circuit QED are often called cavities, a term that was borrowed from cavity QED. These circuit elements behave as harmonic oscillators and are traditionally used to read out qubit states without probing and, thus, disturbing the quantum information directly. In this project, however, resonators are considered for a different purpose: to generate a new kind of qubit more robust to environmental noise. Resonators can be considered protectors of quantum information in integrated circuits. Though they typically have been used to protect quantum information from destruction, we explore how they may also further protect quantum information from noise.

When it comes to microwave engineering and superconducting quantum circuits, there are two common types of resonators used: lumped element resonators and coplanar waveguides. The former will be the focus of this project, though exploration into the latter may prove interesting for future work and should fit within the framework of the methods presented.

A lumped element resonator is comprised of an inductor, L , and a capacitor, C , in parallel (see figure 4), forming a simple electrical LC harmonic oscillator. The name "lumped element" refers to the fact that the physical size is much smaller than the associated wavelength and that the voltage and current are relatively steady across the element [11]. The Lagrangian of such a resonator is

$$\mathcal{L} = \frac{LI^2}{2} - \frac{q^2}{2C} \quad (12)$$

where q is the charge of the capacitor and I is the inductor's current. Recalling from charge conservation that $I = \dot{q}$, which means the Lagrangian may be rewritten as

$$\mathcal{L} = \frac{L}{2}\dot{q}^2 - \frac{1}{2C}q^2 \quad (13)$$

Now this looks remarkably similar to the classical problem of a mass on a spring; here, instead of distance the degree of freedom is charge, rather than mass there is inductance, and the inverse capacitance replaces the spring constant [10]. This similarity is not only incredibly helpful for further analysis of resonators, but will come into play again in section 3.3. Following the Euler-Lagrange equation of motion $\frac{\partial \mathcal{L}}{\partial q} - \frac{d}{dt} \frac{\partial \mathcal{L}}{\partial \dot{q}} = 0$ equation 13 becomes

$$\ddot{q} = -\frac{1}{LC}q = -\omega_r^2 q \quad (14)$$

where $\omega_r = \frac{1}{\sqrt{LC}}$ is the natural frequency of the resonator [10]. $\frac{\partial \mathcal{L}}{\partial \dot{q}}$ defines the conjugate or generalized momentum [12], which, as seen in equation 13, is $L\dot{q}$ or LI , also known as the inductor's flux, Φ [10]. The resonator's Hamiltonian is then expressed as

$$H = \Phi\dot{q} - \mathcal{L} = \frac{\Phi^2}{2L} + \frac{1}{2C}q^2 \quad (15)$$

If, however, the flux of the node next to the capacitor, ϕ , see Figure 4, is used as the degree of freedom instead of charge and the negative charge, $Q = -q$,

becomes the conjugate momentum, then the Hamiltonian may be expressed as

$$H = Q\dot{\phi} - \mathcal{L} = \frac{1}{2C}Q^2 + \frac{\phi^2}{2L} \quad (16)$$

In either case, by following the canonical commutation relation, the degree of freedom and the conjugate momentum may be expressed as operators that are Fourier transforms of each other

$$[\hat{q}, \hat{\Phi}] = i\hbar = [\hat{\phi}, \hat{Q}] \quad (17)$$

Now equation 16 can be expressed in terms of raising and lowering operators [10].

$$\begin{aligned} \hat{a} &= i\frac{1}{\sqrt{2C\hbar\omega_r}}\hat{Q} + \frac{1}{2L\hbar\omega_r}\hat{\phi} \\ \hat{a}^\dagger &= -i\frac{1}{\sqrt{2C\hbar\omega_r}}\hat{Q} + \frac{1}{2L\hbar\omega_r}\hat{\phi} \\ [\hat{a}, \hat{a}^\dagger] &= 1 \end{aligned} \quad (18)$$

Now the Hamiltonian of the resonator can be written as [10]

$$H = \frac{\hbar\omega_r}{2}(a^\dagger a + aa^\dagger) = \hbar\omega_r(a^\dagger a + \frac{1}{2}) \quad (19)$$

Expressing the resonator Hamiltonian in terms of raising and lowering operators will come up again when the Jaynes-Cummings Hamiltonian is discussed, which explains the interactions between qubits and resonators.

LC resonators are often referred to as quantum harmonic oscillators as they can be expressed by an evenly spaced (harmonic) potential energy well. This means that the energy levels of a resonator are evenly spaced by $\hbar\omega_r$. Qubits, however, should only have two energy levels, the ground and first excited state, as was discussed in the introduction; therefore, transmons – a commonly used type of superconducting qubit – are anharmonic quantum oscillators, meaning they have unevenly spaced energy levels. This allows access to the first two energy levels in isolation from the rest (see figure 5). Next, a deeper understanding of the transmon is presented.

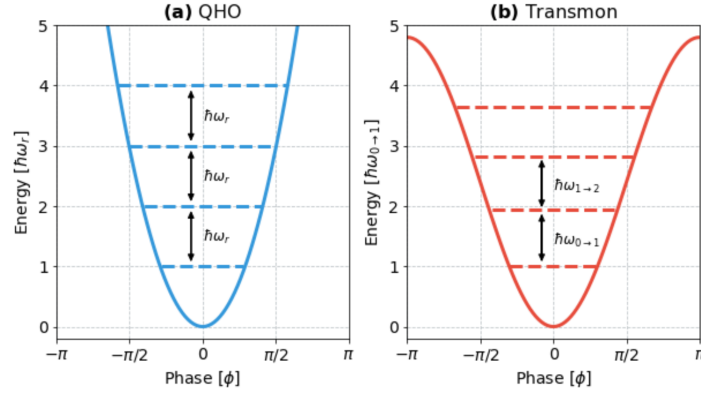


Figure 5: The energy potentials for a quantum harmonic oscillator (resonator) and transmon. The former shows energy levels evenly spaced by $\hbar\omega_r$, where ω_r is the frequency of the resonator. The transmon energy levels are unevenly spaced. They depend on the frequency associated with transitioning between each respective energy level, e.g. the gap between the first two levels is proportional to $\omega_{0 \rightarrow 1}$, the frequency of the transmon transitioning from the ground to the first excited state. Other spacings between energy levels are not proportional to this same frequency, though. Image is from Hoffer (2021) [13].

2.3 Transmon and Jaynes-Cummings Hamiltonian

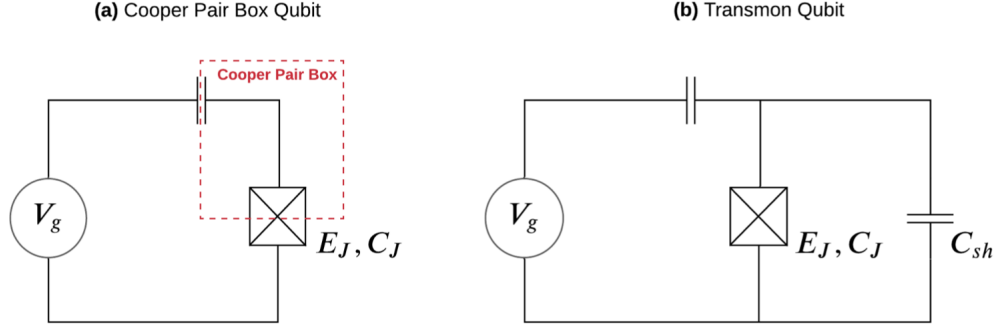


Figure 6: The Cooper pair box circuit (left) is the same as the LC resonator, but with the inductor of the resonator replaced by a Josephson junction, which is marked by the box with an x in it. The transmon circuit (right) is similar to the Cooper pair box, but includes a shunt capacitor, C_{sh} , to minimize charge noise. E_J and C_J refer to the Josephson energy and capacitance respectively [13, 14]. Image is from Hoffer (2021) [13].

Superconducting qubits grew out of a simple arrangement referred to as the Cooper pair box, which looks like the LC resonator already discussed, but with a JJ replacing the inductor (see figure 6). “The Josephson junction is the only known dissipationless non-linear circuit element” [14]. As previously discussed, if one considers the simple LC resonator, as introduced in section 2.2, with a Josephson junction replacing the inductor, then the harmonic oscillations become non-linear, meaning that its associated energy levels are no longer evenly spaced. Thus, if the Cooper pair box is engineered with specific capacitance and inductance, then the energy levels can be designed in such a way where the first two states are within frequency range of experimental interest (gigahertz) and the higher energy bands are out of range [14]. With the first two non-linear energy levels – the ground state, $|0\rangle$, and first excited state, $|1\rangle$ – accessible, this circuit element now behaves as a physical qubit when cooled down to roughly 10mK and manipulated with microwave pulses.¹

The transmon is a widely used updated version of the Cooper pair box in that it is less sensitive to charge noise compared to its predecessor [14]. It is characterized by the ratio of its Josephson energy, E_J , to its charging energy, E_C . The former relates to when a pair of spin-1/2 particles tunnel across the JJ. This tunneling generates “charge difference” between the two capacitors of

¹Please note that the word “physical” is consciously included as often times when quantum computers are discussed in terms of use cases, the size or number of qubits is described in terms of logical qubits. A single logical qubit is made of several physical qubits along with some error correcting code to mitigate natural, unavoidable errors that arise in the system. As this project is considering hardware and not quantum algorithms, when qubits are mentioned the implication is that they are physical qubits, not logical ones.

the JJ, thus determining the qubit state [14]. Charging energy, on the other hand, relates to the total capacitance of the transmon $E_C = \frac{e^2}{C_{sh} + C_J}$ [13, 14]. Transmons have a ratio of Josephson energy to charging energy such that $\frac{E_J}{E_C} > 4$ [14]. Increasing E_J/E_C decreases tunneling and the charge noise across the junction [14]. For this reason, some researchers use $\frac{E_J}{E_C} = 50$ [13]. The bare qubit frequency – i.e. the frequency of a single qubit in isolation, not coupled to any other qubits or resonators – is also dependent on the Josephson and charge energies, $\omega_q \approx \sqrt{8E_J E_C}$ [13].

Josephson and charge energy also help describe a transmon's rate of interaction with a resonator, often called the coupling, g [14].

$$g_{i,j} = \sqrt{2}g\left(\frac{E_J}{8E_C}\right)^{1/4} \langle i | (b - b^\dagger) | j \rangle \quad (20)$$

b and $b^\dagger$ are the raising and lowering operators of the qubit, while $|i\rangle$ and $|j\rangle$ refer to the two qubit states of interest. If these two states are nearest neighbors, then this expression simplifies to the following [14].

$$g_{j,j+1} = g(\sqrt{2(j+1)})\left(\frac{E_J}{8E_C}\right)^{1/4} \quad (21)$$

Though it is nice to express the coupling between a qubit and resonator, what is measured is actually another parameter called the dispersive shift, χ , which is expressed in terms of the coupling, as well as the difference between qubit and resonator frequencies, $\Delta = \omega_{ij} - \omega_r$. ω_{ij} is the frequency of a qubit shifting between states $|i\rangle$ and $|j\rangle$. Δ is often called the detuning. As qubits cannot be measured directly without destroying their quantumness, they are attached to resonators. A resonator is then measured and if its frequency has shifted slightly, it is indication that it is attached to a qubit. This shift marks not only the presence of a qubit, but also which state the qubit is in and how strongly coupled the resonator and qubit are. This shift is the dispersive shift. If the resonator's frequency decreases by χ then the qubit is in the ground state whereas if it increases by χ , then the qubit is in the excited state.

$$\chi_{ij} \equiv \frac{g_{ij}^2}{\Delta_{ij}} \quad (22)$$

Using equation 20, then this becomes

$$\chi_{ij} = 2\sqrt{\frac{E_J}{8E_C}} \frac{g^2 |\langle i | (b - b^\dagger) | j \rangle|^2}{\Delta_{ij}} \quad (23)$$

To further simplify things, an effective dispersive shift can be defined as follows [14].

$$\chi_{eff} = \chi_{01} + \frac{\chi_{12} + \chi_{02}}{2} \propto \frac{g^2}{\Delta} \quad (24)$$

The effective dispersive shift is useful for properly expressing the effective Hamiltonian of the single qubit-single resonator system.

$$H_{eff} = \frac{\hbar(\omega_{01} + \chi_{01})}{2} + \hbar((\omega_r - \chi_{01} - \chi_{02}) + \chi_{eff}\sigma_z)a^\dagger a \quad (25)$$

This expresses the state of the qubit in the first term and the resonator in the second term with a and $a^\dagger$ as the lowering and raising operators associated with the photon or resonator, as was noted in section 2.2. The $\chi_{eff}\sigma_z$ term represents the interaction between the two elements. This Hamiltonian is commonly referred to as the Jaynes-Cummings Hamiltonian and is visualized using a "ladder" as seen in figure 7 [14]. The Jaynes-Cummings Hamiltonian derivation is a thorough method for understanding the dynamics of a transmon-resonator system and may very well be useful for future generations of the two-dimensional LC resonator lattice; however, this thesis implemented a more general qubit-resonator coupling scheme that is not necessarily within the dispersive regime introduced by the derivation above (see section 5.1).

Superconducting qubit-resonator systems are excited and controlled using microwave pulses. That is how qubit gates are run and, consequently, how qubit states are changed. In this project the new proposed qubit is actually a bunch of resonators and we study how to control this array using transmons. In this way microwave pulses may excite the array's edge mode, which subsequently could excite the attached qubit(s), which in turn could change the chirality of the edge modes. Modeling this proposed process requires an understanding of edge modes and, thus, properties often associated with topological insulators.

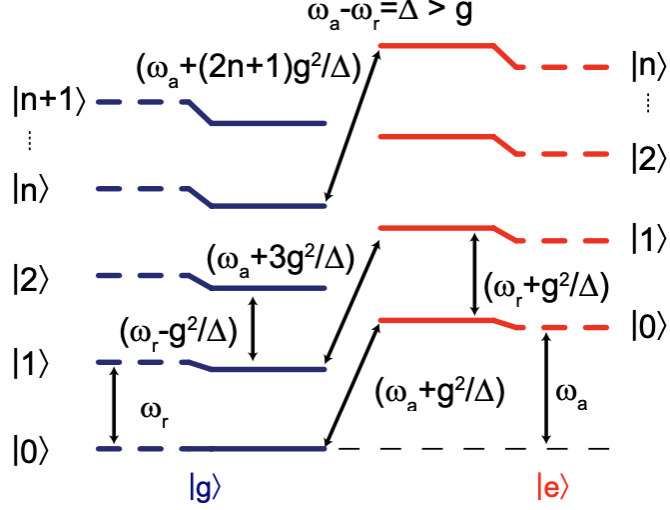


Figure 7: Interactions described by the Jaynes-Cummings Hamiltonian of a single qubit and resonator pair can be represented in a ladder schematic. The blue lines on the left represent when a qubit with bare frequency ω_a is in the ground state, $|g\rangle$, and the red lines are when it is in the excited state, $|e\rangle$. The dashed lines represent the uncoupled Hamiltonian's eigenstates, whereas the solid lines are the energies when considering coupling. $|n\rangle$ is the photon number. Note that a true qubit would only consider the first two states, $|0\rangle$ and $|1\rangle$, but higher order levels are included for completeness. This image considers the situation known as the dispersive limit, where $\omega_a - \omega_r = \Delta > g$; g is the coupling between qubit and resonator. In this scenario, the "effective" resonator and qubit frequencies shift. The bare qubit and resonator frequencies are shown by both ω_a and ω_r marked by the lower most vertical arrows between $|0\rangle$ and $|1\rangle$ depicting the frequencies in terms of the uncoupled states. The shifted frequencies associated with the coupled system are noted in terms of coupling and detuning, $\frac{g^2}{\Delta}$. This also refers to the dispersive shift, χ . These shifted frequencies are marked by the arrows relating the solid lines. When the qubit is in the ground state, the resonator's frequency decreases slightly (as is seen in the shift between the dashed and solid blue lines at $|1\rangle$). In contrast, when the qubit is in the excited state, the resonator's frequency increases (as is seen in the shift between the dashed and solid red lines at $|1\rangle$). Image is from Schuster (2007) [14].

3 Topological quantum systems

Although this project focuses on the behavior of quantum circuits, it is motivated by properties commonly found in topological systems, particularly the appearance and robustness of edge modes along a lattice. While traditionally these edge modes appear in quantum systems, they have recently also been demonstrated on an array of classical mechanical oscillators that displayed two modes in the clockwise and counter-clockwise directions [6]. As LC resonators are the electrical equivalent to mechanical oscillators, we believe achieving such edge states on a two-dimensional array of resonators is also possible. Topologically non-trivial systems can host edge modes. For a system to be topological, there must be an energy band gap between what is considered the physical bulk of the system versus its edges and symmetry projection must be present. If both of these conditions are upheld, the system will be exceptionally robust to perturbations, or rather any noise occurring in the environment.

Recent encouraging research has introduced how two one-dimensional chains of LC resonators connected by a qubit can be entangled while preserving topological edge modes on each chain [15]. Such a one-dimensional chain is explained by the Su-Schrieffer-Heeger (SSH) model. Though the proposed work of this thesis is not in one-dimension and, thus, cannot directly utilize the SSH model, the SSH model is a simpler, perhaps even more intuitive presentation of the importance of energy band gap and chiral symmetry. A further analysis of topological systems continues with a focus on a more directly relevant model: the quantum spin Hall effect (QSHE) [16]. For this, a quick introduction to time-reversal symmetry and Kramers pairs is provided. Finally, a review of Süsstrunk and Hubers’ (2015) work demonstrating chiral edge modes on mechanical oscillators is presented [6]. This serves as a guidance for how to set up the Hamiltonian of our proposed two-dimensional LC resonator system.

3.1 SSH Model

The SSH model describes the topological behavior of a one-dimensional chain of particles or atoms. A common physical example for which the SSH model applies is polyacetylene [18]. Polyacetylene is a chain of carbon atoms, each with one hydrogen atom hanging off. When the chain is broken up into sites – each with one unique pair of carbon atoms – then two types of general bonds can be identified: inter- and intra-site bonds, whose strengths are denoted w and v respectively [17, 19]. These are shown as a single (w) and double (v) line connecting the carbon atoms in figure 8. Thus, a single site or unit cell is comprised of one atom on the bottom connected to one atom on top by a double line (v). All “bottom” (“top”) particles comprise what is considered a sublattice. Each sublattice’s particles do not interact with one another. A particle can only interact with its nearest neighbor, which is of opposite sublattice type (i.e. top only interacts with its nearest neighbor bottom particles and vice versa). This sublattice symmetry is the SSH model’s chiral symmetry [17, 19]. It is useful to note that this scenario considers non-interacting electrons, thus each electron

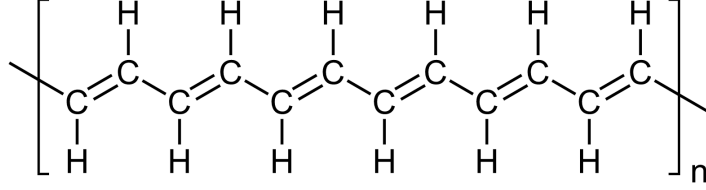


Figure 8: A chain of carbon atoms known as polyacetylene is a common physical example of the SSH model. Atoms are categorized into sub-lattices. This is demonstrated by the staggering of the lattice showing half the particles on "top" and half on "bottom." One unit cell is comprised of one particle from each sublattice, i.e. one "top" and one "bottom." The bond of two atoms within a cell is denoted by a double line and is referred to as v . The intercell bonds are shown as a single line and called w . v and w are referred to as hopping amplitudes [17].

can be considered a single particle [19]. The Hamiltonian for such a system is given as

$$\begin{aligned}
 H = & v \sum_{m=1}^N |m, B\rangle \langle m, A| + |m, A\rangle \langle m, B| \\
 & + w \sum_{m=1}^{N-1} |m+1, B\rangle \langle m, A| + |m, A\rangle \langle m+1, B|
 \end{aligned} \tag{26}$$

where N is the number of sites in the chain and A and B (what was previously referred to as top and bottom) denote which sublattice, as seen in Figure 8 [19].

The Pauli exclusion principle requires that no two fermions, or spin-half particles, which includes electrons, exist in the same state at the same time in the same system [20]. Given no potential energy and zero temperature, this means all eigenstates of the SSH Hamiltonian are limited to single occupancy as each site has one particle of each spin; this is also known as half-filling [19]. Expressing the Hamiltonian in terms of internal and external degrees of freedom will become useful as the system is described in terms of bulk and edges. The internal degree of freedom indexes the unit cells and the external degree of freedom identifies the sublattices [19]. These are m and α respectively, such that $|m, \alpha\rangle = |m\rangle \otimes |\alpha\rangle \in \mathcal{H}_{\text{internal}} \otimes \mathcal{H}_{\text{external}}$ [19]. Now the SSH Hamiltonian can be written as

$$\begin{aligned}
 H = & v \sum_{m=1}^N |m\rangle \langle m| \otimes \hat{\sigma}_x \\
 & + w \sum_{m=1}^{N-1} |m+1\rangle \langle m| \otimes \frac{\hat{\sigma}_x + i\hat{\sigma}_y}{2} + |m\rangle \langle m+1| \otimes \frac{i\hat{\sigma}_x - \hat{\sigma}_y}{2}
 \end{aligned} \tag{27}$$

where $\hat{\sigma}_x$ and $\hat{\sigma}_y$ are the Pauli matrices, with $\hat{\sigma}_x$ defined in Equation 5 and

$$\hat{\sigma}_y = \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} \quad (28)$$

In this expression, the hopping amplitudes look like operators on their respective bonds [19].

The Hamiltonian for the bulk should be irrespective of the edges, as the idea that we can distinguish clearly between these two modes is central to topological insulators. The bulk is thus given periodic boundary conditions, letting us effectively treat it as a ring [19]. The bulk Hamiltonian now is

$$\begin{aligned} H_{bulk} = & \sum_{m=1}^N v(|m, B\rangle \langle m, A| + |m, A\rangle \langle m, B|) \\ & + w(|m \bmod N + 1, A\rangle \langle m, B| + |m, B\rangle \langle m \bmod N + 1, A|) \end{aligned} \quad (29)$$

where

$$\begin{aligned} H_{bulk} |\Psi_n(k)\rangle &= E_n |\Psi_n(k)\rangle, \\ n &\in 1, \dots, 2N \end{aligned} \quad (30)$$

with wavenumber or, as is more commonly referred to in condensed matter theory, crystal momentum k [17, 19]. As the bulk is translationally invariant, Bloch's theorem allows expressing the state in terms of a plane wave in the first Brillouin zone, $k \in \delta_k, \dots, \delta_k N$ where $\delta_k = \frac{2\pi}{N}$,

$$|k\rangle = \frac{1}{\sqrt{N}} \sum_{m=1}^N e^{imk} |m\rangle \quad (31)$$

with Bloch eigenstates $|\Psi_n(k)\rangle$

$$|\Psi_n(k)\rangle = |k\rangle \otimes |u_n(k)\rangle \quad (32)$$

$$|u_n(k)\rangle = a_n(k) |A\rangle + b_n(k) |B\rangle \in \mathcal{H}_{int} \quad (33)$$

$|u_n(k)\rangle$ are the eigenstates of the bulk momentum-space Hamiltonian

$$H(k) |u_n(k)\rangle = E_n(k) |u_n(k)\rangle \quad (34)$$

where

$$\begin{aligned} H(k) &= \langle k| H_{bulk} |k\rangle \\ &= \sum_{\alpha, \beta \in A, B} \langle k, \alpha| H_{bulk} |k, \beta\rangle \cdot |\alpha\rangle \langle \beta| \end{aligned} \quad (35)$$

[19].

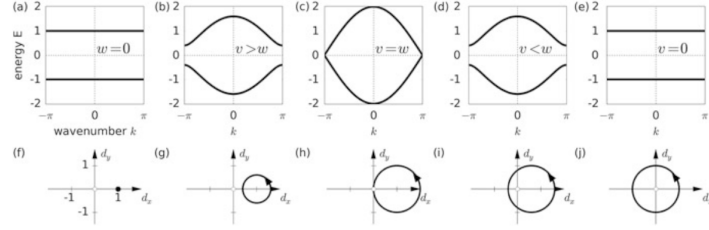


Figure 9: (a) to (e) show the energy as a function of wavenumber for varying relations between inter- and intra-cell hopping amplitudes, w and v : (a) $v = 1, w = 0$, (b) $v = 1, w = 0.6$, (c) $v = 1, w = 1$, (d) $v = 0.6, w = 1$, (e) $v = 0, w = 1$ [19]. The smallest gap between the energy bands is equal to twice the absolute value of the difference between the hopping amplitudes. Note that (c) shows a conductor as the energy bands connect allowing electrical flow without tunneling. In the insulator case, as the gap increases the number of occupied energy states decreases. (f) to (j) show the circular $\mathbf{d}(k)$ path for the bulk momentum-space Hamiltonian as k crosses through the Brillouin zone from 0 to 2π for each respective situation shown above in (a) to (e) [19]. Image is from Asbóth, Oroszlány, Pályi (2016) [19].

The function $u_{n,k}(x)$ (where $\Psi_{n,k}(x) = e^{ikx}u_{n,k}(x)$) demonstrates periodicity in real space, but not in the Brillouin zone; however, the Fourier transform from real space to Brillouin zone is on the external degree of freedom only, thus periodicity in the Brillouin zone is maintained [19].

$$\begin{aligned} H(k + 2\pi) &= H(k) \\ |u_n(k + 2\pi)\rangle &= |u_n(k)\rangle \end{aligned} \quad (36)$$

The bulk momentum-space Hamiltonian has a zero diagonal. This reflects that there is no interaction among particles within a given sublattice. Particles in an SSH chain only interact with their nearest neighbors [19]. This localized interaction denotes a chiral symmetry in the system. The Hamiltonian can then be written as

$$H(k) = \begin{pmatrix} 0 & v + we^{-ik} \\ v + we^{ik} & 0 \end{pmatrix} \quad (37)$$

Given Equation 34, $H(k)^2 = E(k)^2 \hat{\mathbb{I}}_2$,

$$E(k) = \pm |v + we^{ik}| = \pm \sqrt{v^2 + w^2 + 2vw \cos(k)} \quad (38)$$

This allows us to see beauty of the SSH model, which arises when the hopping amplitudes v and w are staggered [19]. To fully appreciate when staggered v and w give rise to a topologically invariant system (as opposed to a trivial one), though, the Hamiltonian can be rewritten in terms of something called the d -vector, $\mathbf{d}(k)$.

Given $|u_n(k)\rangle \in \mathcal{H}_{internal}$ provides information on the internal structure, the bulk momentum-space Hamiltonian for a system with two states per cite (i.e. two particles per unit cell) can be expressed as

$$\begin{aligned} H(k) &= d_0(k)\hat{\sigma}_0 + d_x(k)\hat{\sigma}_x + d_y(k)\hat{\sigma}_y + d_z(k)\hat{\sigma}_z \\ &= d_0(k)\hat{\sigma}_0 + \mathbf{d}(k)\hat{\sigma} \end{aligned} \quad (39)$$

where

$$\begin{aligned} d_0(k) &= 0 \\ d_x(k) &= v + w \cos k \\ d_y(k) &= w \sin k \\ d_z(k) &= 0 \end{aligned} \quad (40)$$

Thus, the direction of $\mathbf{d}(k)$ determines the internal structure of eigenstates and the magnitude of the vector determines the energy. The radius of the circle is of size w and the center sits at v . $\mathbf{d}(k)$ must be a closed loop and encircle the origin for the lattice to maintain chiral symmetry and be considered topological [19]. This becomes clear in figure 9, in which the top row shows energy plotted against wavenumber for varying values of v and w and the bottom row plots the associated $\mathbf{d}(k)$ to the image above it. The image to the far right (e) reflects a system in which there is no intra-cell hopping, i.e. the two particles within a unit cell are fully disconnected. In this case it is clear how there would be a separation in what is considered the bulk compared to the edges as there is a physical break in the connections. The next image to the left (d) demonstrates a system in which the intra-cell hopping has been slightly turned on, but is still weaker than inter-cell hopping. The particles that were physically separated in (e) on the edges are now connected again; however, the energy band gap is still present and $\mathbf{d}(k)$ still encircles the origin. As both chiral symmetry and the energy band gap have been maintained, this still reflects a topological system. The central image (c), however, shows a system in which $v = w$ and, thus, there is closure to the energy band gap. This gap closure demonstrates the topological phase transition between $v < w$ and $v > w$ or rather between the system behaving topologically or trivially. The following systems (a) and (b) restore the band gap, but $\mathbf{d}(k)$ no longer contains the origin and, therefore, are considered trivial rather than topological; they no longer display separation in bulk and edge modes [19]. By mapping the energies for various combinations of hopping amplitudes and $\mathbf{d}(k)$ of the bulk momentum-space Hamiltonian, we begin to see the significance of bulk versus edge modes.

Figure 10 further emphasizes how when the intra-cell hopping amplitude, v , is turned on and increased, the system still remains topologically invariant until it equals w [19]. This is visible in figure 10(a), in which there is a line around $E = 0$, which is separate from the rest of the energy spectrum. This line reflects the energy of the edge mode. It starts at 0 when $v = 0$ and increases exponentially as v increases. Recall that when $v = 0$ the edges of the lattice are physically separated from the rest of the system, hence zero energy. The delayed union of the edge and bulk modes maintains a gap between the two energies

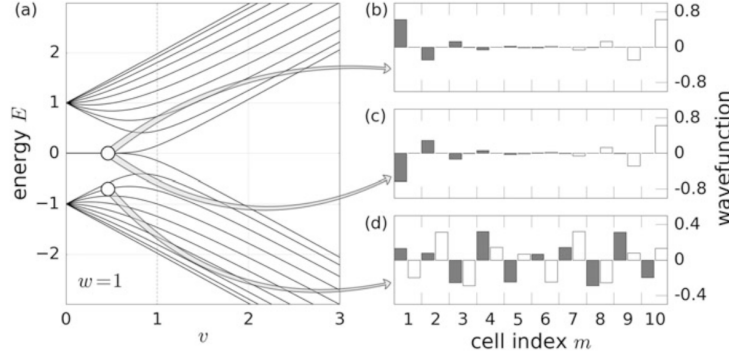


Figure 10: (a) The energy mapped as a function of intra-cell hopping, v , of a ten-cell lattice described by the SSH model with inter-cell hopping $w = 1$. When $v = 0$, the cells along the chain are broken apart, thus creating physically separated edges. As v increases, however, the energy of the edge – denoted by the seemingly horizontal line at $E = 0$ – exponentially increases. Until $v = 1$, the edge cells – though physically connected again – still maintain their own distinguishable energy separate from the bulk of the lattice. After $v = 1$, the distinct edge modes disappear and the system is no longer topological. (b) - (c) These plots highlight how a wavefunction associated with the edge mode in (a) are indeed isolated to the edge cells. (d) When a wavefunction is considered along the bulk of the system, it is spread amongst all cells and not isolated in the same way it is for edges as shown in (b) and (c) [19]. Image is from Asbóth, Oroszlány, Pályi (2016) [19].

even as v increases (but is still less than w). This allows the system to remain topologically nontrivial. Figure 10(b) and (c) highlight how a wavefunction associated with the zero energy edge modes of the topological system indeed are isolated to the edges. Whereas figure 10(d) demonstrates how a wavefunction associated with the bulk states is spread across the entire bulk and not isolated to any particular cells in the chain. Once $v = w$, however, and as it continues to increase, the system loses its topological classification and becomes trivial; the edges are indistinguishable energetically from the bulk [19].

After examining this one-dimensional example of the SSH model, it is clear how edge modes can exist in topological systems. It is also clear how maintaining an energy band gap and symmetry is crucial to maintaining topological invariance [19]. Where the SSH model represents these features in terms of hopping amplitudes that have simple-to-understand physical meaning (e.g. the bond strength between two carbon atoms in polyacetylene), other symmetries in higher dimensional systems may not be as obvious at first. One of them, time-reversal symmetry, is certainly not as intuitive as bond strengths, but is critical for demonstrating the QSHE, which describes how spin-up and spin-down particles can exist in the same system at the same time.

3.2 Time-Reversal Symmetry and Quantum Spin Hall Effect

Symmetry operators, which are unitary, do not disturb the inner product of two states

$$\langle U\psi|U\phi\rangle = \langle\psi|U^\dagger U|\phi\rangle = \langle\psi|\phi\rangle \quad (41)$$

and, thus,

$$U^\dagger H U = H \quad (42)$$

where H is a Hamiltonian with symmetry U [21]. Anti-unitary symmetry operators, Ω , however, do not preserve inner products and instead produce the complex conjugate [21].

$$\langle\Omega\psi|\Omega\phi\rangle = \langle\phi|\psi\rangle = \langle\psi|\phi\rangle^* \quad (43)$$

Similarly, where symmetry operators are linear $U(\alpha|\phi\rangle + \beta|\psi\rangle) = \alpha U|\phi\rangle + \beta U|\psi\rangle$, anti-unitary operators are anti-linear

$$\Omega\alpha|\phi\rangle = \Omega\alpha\Omega^{-1}\Omega|\phi\rangle = \alpha^*\Omega|\phi\rangle = \alpha^*|\Omega\phi\rangle \quad (44)$$

where $\alpha \in \mathbb{C}$ [21]. Rather, an anti-unitary operator transforms a complex number into its complex conjugate, which agrees with the same behavior seen in relation to an inner product. Given this, we can say that $\Omega = \mathcal{K}$, where $\mathcal{K}$ is the complex conjugation operator with $\mathcal{K}^2 = \mathbb{I}$ and $\mathcal{K} = \mathcal{K}^{-1}$ [21]. To see how this affects quantum mechanical systems, the simplest thing to do is to consider $\mathcal{K}$ on one-dimensional Schrödinger's equation, $i\hbar \frac{\partial\psi(x,t)}{\partial t} = H\psi(x,t)$

$$Ki\hbar \frac{\partial\psi(x,t)}{\partial t} K = K H K \cdot K\psi(x,t) K \quad (45)$$

$$-i\hbar \frac{\partial\psi^*(x,t)}{\partial t} = H^*\psi^*(x,t) \quad (46)$$

$$i\hbar \frac{\partial\psi^*(x,t)}{\partial -t} = H^*\psi^*(x,t) \quad (47)$$

With $H = H^*$, then $\psi^*(x,t)$ is the solution to the time-reversed one-dimensional Schrödinger's equation [21]. In a system with spin, the Hamiltonian to describe spin orbit coupling is dependent on the Pauli matrices; however, σ_y is not $\mathcal{K}$ invariant [21]. Therefore, to maintain time-invariance for spin systems, another anti-unitary operator, $\mathcal{T}$, is defined [6, 21].

$$\mathcal{T} = i\sigma_y \mathcal{K} \quad (48)$$

When considering how $\mathcal{T}$ works with eigenstates, $\mathcal{T}|\psi\rangle = \lambda|\psi\rangle$, it becomes evident that it cannot have any [21]. This is seen by adding an extra time-reversal operator, $\mathcal{T}^2|\psi\rangle = \mathcal{T}\lambda|\psi\rangle$. By definition, the left-hand side should equal $-|\psi\rangle$. The right-hand side, however, does not [21].

$$\mathcal{T}^2|\psi\rangle = \mathcal{T}\lambda|\psi\rangle = \lambda^*\mathcal{T}|\psi\rangle = |\lambda|^2|\psi\rangle \neq -|\psi\rangle \quad (49)$$

In this way there are two degenerate eigenstates. $|\psi\rangle$ and $|\phi\rangle$ are referred to as a Kramers pair [21].

$$\begin{aligned}\mathcal{T}|\phi\rangle &= -|\psi\rangle \\ \mathcal{T}|\psi\rangle &= |\phi\rangle\end{aligned}\tag{50}$$

These states are orthogonal as

$$\langle\psi|\mathcal{T}\psi\rangle = \langle\mathcal{T}^2\psi|\mathcal{T}\psi\rangle = -\langle\psi|\mathcal{T}\psi\rangle\tag{51}$$

Additionally, the time-reversal symmetry operator does not have a Hermitian conjugate [21].

$$\langle\psi|\mathcal{T}\phi\rangle \neq \langle\mathcal{T}^\dagger\psi|\phi\rangle\tag{52}$$

This equation would require that the left and right hand sides demonstrate anti-linear and linear relationships, respectively [21].

Applying time-reversal symmetry to a Hamiltonian describing a system with spin demonstrates the existence of a Kramers pair; one eigenstate may be spin up with the other is spin down and while they exist at the same energy, they do not interact. This phenomenon is known as the quantum spin Hall effect [16]. QSHE demonstrates topological systems through time-reversal symmetry. Projecting this concept onto a system of LC resonators, there are two chiral edge modes – clockwise and counter-clockwise – that exist at the same frequency but do not cancel each other out. This application of QSHE in classical oscillators was first demonstrated on a lattice of mechanical pendula by Süsstrunk and Huber (2015) [6]. To understand what the Hamiltonian looks like for a two-dimensional array of LC resonators mimicking the QSHE, it is useful to understand the Hamiltonian for such a mechanical system.

3.3 Quantum Spin Hall Effect in the Classical Realm

The motivation behind finding topological edge modes in LC resonators came from results showing such behavior in a two-dimensional array of pendula coupled by springs [6]. Neither of these systems are ruled by the laws of quantum mechanics, yet given that the electrical oscillators behave in a similar manner to the mechanical ones (see section 2.2), it is believed that such methods can be extended to electrical oscillators. The overview of how such mechanical systems are compared to quantum-associated topological insulators is as follows: QSHE reflects a separation between edge and bulk modes in a system with two degrees of freedom reflecting two spin modes. Classical oscillators naturally only describe systems with one degree of freedom. To remedy this, consider the Hamiltonian describing the classical oscillator system, H_ϕ ; when the complex conjugate of the original Hamiltonian, H_ϕ^* , is also included, then requirements for comparing to QSHE are met as there are now two degrees of freedom contained in the whole system's Hamiltonian, H [6, 19, 22]. This is illustrated in equation 53. H_ϕ is taken from Hofstadter's model (see equation 54), which is used in condensed matter theory to describe flux under a very strong magnetic

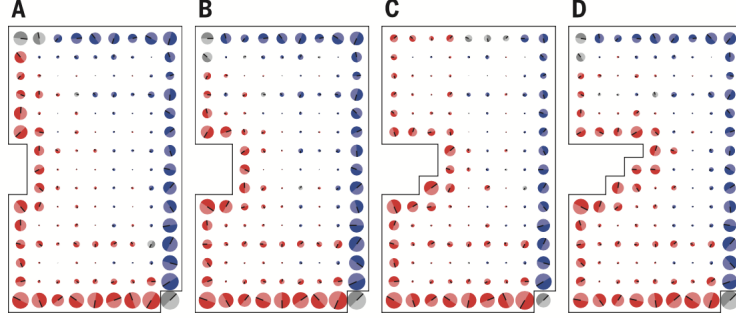


Figure 11: Four deformations of a two-dimensional lattice of pendula are shown. Each deformation of the physical edge is considered a perturbation to the system. The system, however, is undisturbed by such environmental changes as it is topologically protected. The red and blue coloring around the edges denote the two circular edge modes observable after exciting the site furthest to the bottom right. Image is from Süsstrunk and Huber (2015) [6].

field. Though there is no exceptional magnetic field in the case of mechanical or LC oscillators, there is still a flux acquired from the phase in both systems that can be described the same way as in Hofstadter's model [6, 23]. As the system is mimicking the QSHE on spin-1/2 particles, the associated flux is $\phi = \frac{2\pi}{3}$ [6]; a single unit cell has flux ϕ and three unit cells make up a flux cell, which collectively has a flux of 2π .

$$\mathcal{H} = \begin{pmatrix} H_\phi & 0 \\ 0 & H_\phi^* \end{pmatrix} \quad (53)$$

$$H_{\alpha,\phi} = f_0 \sum_{r,s,\alpha} |r, s, \alpha\rangle \langle r, s \pm 1, \alpha| + |r, s, \alpha\rangle \langle r \pm 1, s, \alpha| e^{\pm i\alpha\phi_s} \quad (54)$$

where α represents what is commonly referred to as the spin of the electron, but here differentiates between H_ϕ and H_ϕ^* . (r, s) designate the location of the cell of interest in the lattice (as is noted in figure 12), f_0 is the hopping amplitude, and $\phi_s = \phi s$ [6]. In this convoluted way, condensed matter theory is used to describe a classical system that exhibits behavior previously expected only of quantum systems.

Now H describes the whole system with two degrees of freedom and it should be possible to see how edge modes arise. The problem, however, is that H is complex, but mechanical oscillators are described by Newton's equation of motion, $\ddot{x}_i = \mathcal{D}_{ij}x_j$, with x the coordinate of the pendula. This expresses a real, not complex system. The Hamiltonian expressing the couplings, $\mathcal{D}$, within the mechanical oscillator array must be real, symmetric, and positive semi-definite. By applying a unitary transformation, U , on H , the classical system is finally described fully as both real and with two degrees of freedom [6]. Equations 55 and 56 demonstrate how this transformation maps the quantum, complex

description of the system (i.e. in terms of Hamiltonian H with $\psi_{r,s}^+$ and $\psi_{r,s}^-$ as the wavefunctions at lattice site (r, s) with spin $+$ or $-$) to the real domain of $\mathcal{D}$ (with $x_{r,s}$ and $y_{r,s}$ representing each pendulum at lattice site (r, s)) [6].

$$U = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & -i \\ 1 & i \end{pmatrix} \quad (55)$$

$$\begin{pmatrix} x_{r,s} \\ y_{r,s} \end{pmatrix} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ i & -i \end{pmatrix} \begin{pmatrix} \psi_{r,s}^+ \\ \psi_{r,s}^- \end{pmatrix} \quad (56)$$

where now

$$U^\dagger \mathcal{H} U = \begin{pmatrix} \text{Re}\mathcal{H}_\phi & \text{Im}\mathcal{H}_\phi \\ -\text{Im}\mathcal{H}_\phi & \text{Re}\mathcal{H}_\phi \end{pmatrix} \equiv \mathcal{D} \quad (57)$$

The diagonal elements of the coupling matrix, $\mathcal{D}$, represent coupling between two pendula, either xx or yy, or in terms of the language discussed in the SSH model, two atoms on the same sublattice type in neighboring unit cells [6]. Please remember that as we are no longer ruled by the SSH model, this type

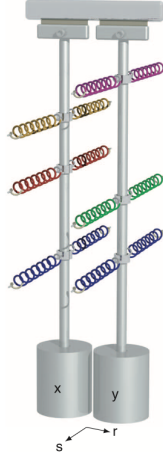


Figure 12: Two mechanical oscillators attached by a series of pairs of springs. The two oscillators create one site or cell in the same way that two carbons create a unit cell in polyacetylene chain. Image is from Ssstrunk and Huber (2015) [6].

of coupling is not forbidden; where the SSH model's chiral symmetry related to sublattice isolation, the QSHE relates to preserving pseudo-spin [6, 19]. The off-diagonal, or imaginary parts, correspond to xy coupling [SH]. Suppose (r,s) is a lattice site. $x_{r,s}$ and $y_{r,s}$ are 1D oscillators in one unit cell at location (r,s) ; therefore, one cell is a two-dimensional oscillator [6]. U allows us to see how the array's eigenmodes are associated with clockwise and counter-clockwise polarized motion, defined as $\alpha = \pm$; "the existence and properties of edge modes are features of the eigenstates of $\mathcal{H}$ or $\mathcal{D}$ " [6]. These chiral modes are to this

classical oscillator problem as pseudo-spin up and down are to the quantum problem presented in the QSHE. What QSHE would consider spin up (down), this problem considers right (left) polarized edge mode [6]. It is $\mathcal{D}$ that will demonstrate QSHE properties through time reversal symmetry and the two chiral edge modes [6, 22].

To understand how time reversal and pseudo-spin symmetries observed in the QSHE for a spin-1/2 quantum state are understood in the case of classical mechanical oscillators, a latitude and longitude are mapped on the Bloch sphere. These are $\varphi = [-\frac{\pi}{2}, \frac{\pi}{2}]$ and $\theta = [0, 2\pi)$ respectively [6]. Such a state can be written as

$$|\psi\rangle = \sin(\frac{\varphi + \frac{\pi}{2}}{2}) |\uparrow\rangle + e^{i\theta} \cos(\frac{\varphi + \frac{\pi}{2}}{2}) |\downarrow\rangle \quad (58)$$

Time reversal moves a point on the Bloch sphere to its anti-nodal point, or a flip along both the latitude and longitude of the sphere $\mathcal{T} : (\varphi, \theta) \rightarrow (-\varphi, \theta + \pi)$, whereas pseudo-spin symmetry only causes a flip along the longitude $\sigma_z : (\varphi, \theta) \rightarrow (\varphi, \theta + \pi)$ [6]. In this way, the mapping of time reversal symmetry in QSHE to symmetries in an array of mechanical oscillators is

$$\mathcal{T} \rightarrow s \circ \mathcal{T} \quad (59)$$

where s represents the symmetry found in local modes of a unit cell, such that $s : (x, y) \rightarrow (y, -x)$ [6].

The assumption Süsstrunk and Huber make is that both s and $\mathcal{T}$ are independently true, or rather that there is no coupling between the left and right circular modes. They do briefly mention the possibility that this is not true. In such an instance $s \circ \mathcal{T}$ is still a symmetry of the system, but s and $\mathcal{T}$ are not symmetries on their own [6]. Süsstrunk and Hubers' consensus is that this latter symmetry scenario is not applicable to pendula [6]. Whether or not it is applicable to an array of LC resonators is something to be determined. What is definitely applicable, though, is the general symmetry analysis comparing classical oscillators to the QSHE, as well as how the coupling matrix or Hamiltonian of the system is generated. Precisely how an array of LC resonators behaves in a similar set up is outlined in chapter 5.

4 Optimal Control

Pulse optimization can take the form of manipulating the pulse's shape (e.g. Gaussian or square), amplitude, and frequency. Determining the optimal control values for the pulse can greatly affect the observed behavior of the system in question and, thus, optimal control is an active area of research [7] [8] [13] [24]. Commonly used methods relate to incorporating the gradient of the pulse into the pulse itself [7] [8] [24]. One of these methods is known as gradient ascent pulse engineering (GRAPE) [7] [24]. This method discretizes the pulse in time, whereas the second method, derivative removal by adiabatic gate (DRAG), maintains a continuous pulse shape, e.g. a Gaussian, and was introduced to minimize qubit information leakage – for example, when a qubit accidentally excites into the $|2\rangle$ state – but could also be used for undesirable couplings between qubits or qubits and resonators [8]. A DRAG-inspired method was implemented, rather than either technique in its pure form. In this chapter GRAPE and DRAG are introduced for the sake of completeness regarding motivation, as well as potential future use; section 5.3 will discuss what pulse optimization method was used.

4.1 GRAPE

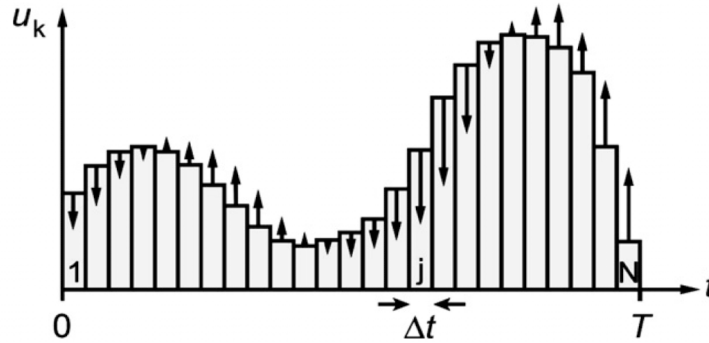


Figure 13: GRAPE algorithm in a single iteration. Each slice in time, j , has a constant control value, u_k , and width Δt . There are N slices total. Each j then has its own gradient calculated, identifying how to adapt the next iteration appropriately. In this way GRAPE addresses all of time in a single iteration. Image is from Khaneja, Reiss, Kehlet, et al (2005) [7].

Traditionally, optimization techniques that utilize a system's gradient use difference method protocols. These methods require discretization of the system into a grid. Consider an example of 500 time steps, N , all of equal size $\Delta t = \frac{T}{N}$, where T is the total time of interest. There are $m = 4$ number of

control variables (e.g. amplitudes / size of amplitude vector/control vector). A difference method would require $500 \times 4 + 1 = 2001$ "full time evolutions" to see the system advance from time $t = 0$ to T [7]. A gradient assessment method that discretizes the system in constant time steps and evaluates all times steps at once, though, as is done in GRAPE, only require two steps: one to propagate time forward from 0 to T and a second to propagate backwards [7]. GRAPE was first introduced as a method of pulse engineering for nuclear magnetic resonance (NMR) [7]. It works better than difference methods because it takes time slices of (pixelates) the control amplitude over all time. Each of these slices is of constant value with respect to the control variable. The gradient of each slice is calculated; the next increment in the algorithm will produce a sum of the respective previous slice value and its gradient [7]. Consider a time-dependent density operator defined as

$$\dot{\rho}(t) = [(\mathcal{H}_0 + \sum_{k=1}^m u_k(t)\mathcal{H}_k), \rho(t)] \quad (60)$$

where $\mathcal{H}_0$ is the free evolution Hamiltonian of the system, $\mathcal{H}_k$ is the k^{th} radio-frequency Hamiltonian, and $u_k(t)$ is the k^{th} entry in the control vector $u(t)$ that is to be optimized such that $\rho(t)$ transitions from $\rho(t) = \rho_0$ to $\rho(T)$ in time T , where $\rho(T)$ is as close as possible to a desired operator, C . Often, and as is introduced in Khaneja, Reiss, Kehlet, et al (2005) [7], $u(t)$ refers to the amplitude of the pulse. If both the density and desired operators are Hermitian, the overlap – or how close $\rho(T)$ is to C – can be expressed as the inner product

$$\langle C | \rho(T) \rangle = \text{tr}(C^\dagger \rho(T)) \quad (61)$$

If these operators are non-Hermitian, however, then the inner product defines a new variable known as the performance index can be introduced, Φ_0 , such that

$$\Phi_0 = \langle C | \rho(T) \rangle \quad (62)$$

The system's time evolution in a given time step, j , is defined by propagator, U , such that

$$U_j = \exp \left\{ -i\Delta t \left(\sum_{k=1}^m u_k(t)\mathcal{H}_k \right) \right\} \quad (63)$$

allows us to express the final density operator as

$$\rho(T) = U_N \dots U_1 \rho_0 U_1^\dagger \dots U_N^\dagger \quad (64)$$

This can then be plugged into the definition for the performance function (equation 62) and together with equation 61, the performance function can be rewritten as

$$\Phi_0 = \langle \lambda_j | \rho_j \rangle \quad (65)$$

where λ_j and ρ_j are respectively the backpropagated desired operator, C , and the density operator, $\rho(t)$, both at time $t = j\Delta t$ [7]. Thanks to "the fact that

the trace of a product is invariant under cyclic permutations of the factors,” [7] these can be defined as

$$\lambda_j = U_{j+1}^\dagger \dots U_N^\dagger C U_N \dots U_{j+1} \quad (66)$$

$$\rho_j = U_j \dots U_1 \rho_0 U_1^\dagger \dots U_j^\dagger \quad (67)$$

The goal of GRAPE is to find the best pulse, or controls $u_k(t)$, that will drive the density operator, $\rho(t)$, to as close to the desired value, C , as possible over time T . In other words, the algorithm should find the ideal controls u_k that maximizes the performance index, Φ_0 [7]. That algorithm is as follows.

The user starts with an initial guess of the controls, $u_k(j)$, for time step j then calculates ρ_j following equation 67 for all $j \leq N$, where N is the total number of time steps. Then the backward propagation is calculated according to equation 66 for all $j \leq N$ (the initial starting point for λ_j is C). The derivative of the performance index with respect to the controls at each time step is calculated according to

$$\frac{\delta \Phi_0}{\delta u_k(j)} = -\langle \lambda_j | i \Delta t [\mathcal{H}_k, \rho_j] \rangle \quad (68)$$

and, thus, the controls may be updated as

$$u_k(j) = u_k(j) + \epsilon \frac{\delta \Phi_0}{\delta u_k(j)} \quad (69)$$

where ϵ is the step size. This process is repeated, starting again with the calculation of equation 67, until Φ_0 surpasses a predetermined threshold [7]. GRAPE’s assessment of the control vector over all time all at once enables great optimization over many control parameters in a computationally cost-effective manner (as was demonstrated earlier with the example of 4 parameters over 500 time slices). While GRAPE can be used for optimizing many systems beyond NMR-specific use cases, a more recent variation of the GRAPE algorithm is specifically useful for protecting qubits from leaking information to other circuit components or higher energy levels. This method, introduced next, is known as derivative removal by adiabatic gate (DRAG).

4.2 DRAG

Quantum systems are known for being imperfect. Qubits – often referred to as two-level systems to acknowledge that they have two energy levels of interest – often leak energy or information into other channels. As the lattice of resonators is meant to behave as a two-level system, it, too, could possibly experience such leakage. While the lattice demonstrates the robustness of topologically-inspired edge modes, thus preventing its signal from leaking into alternative nearby circuit elements, the signal in one lattice mode (e.g. clockwise) can still leak into the other mode (e.g. counterclockwise). Initial simulations show that this happens when the lattice is connected to a qubit. After the qubit gets excited, there

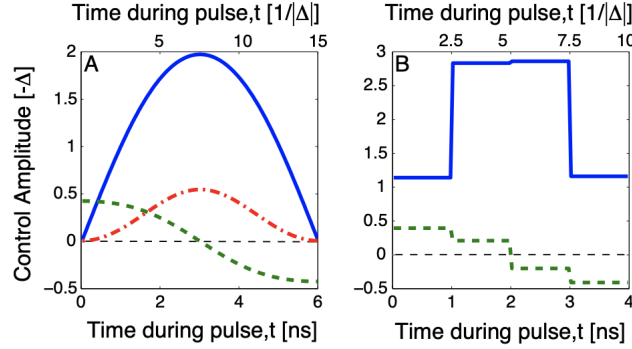


Figure 14: Right: DRAG. Left: GRAPE. Blue lines show $\mathcal{E}^x$ control and green lines are the $\mathcal{E}^y$ control. The red line shows the detuning between the frequencies of the qubit's first excited state and the drive. DRAG is continuous and instantaneous, whereas GRAPE is pixelated (here by 1 ns) and may require multiple iterations to compensate for leakage. Also important to note is that a DRAG pulse always starts and stops at 0. Image is from Motzoi, Gambetta, Rebentros, et al (2009) [8].

is a backpropagation of the signal into the opposite chiral edge mode. While GRAPE could be useful to fix this, there is another quantum optimal control method known as the derivative removal by adiabatic gate (DRAG) [8]. An adiabatic transformation is one that describes a system changing slowly from one Hamiltonian to another; or rather the system is first described by a more easily understood Hamiltonian, then is slowly transformed into the desired Hamiltonian [19, 25]. DRAG pulses are known for minimizing the effect of leakage into higher energy states of qubits, thus it is a good candidate for improving the lattice's response after an attached qubit has been excited [8, 26]. It was also originally suggested that this method may be useful in "turning off unwanted coupling in multi-qubit and qubit-oscillator systems for single and two-qubit operations" [8].

The essential process of DRAG is to simply add the derivative of the initial pulse to the pulse itself. More thoroughly, introducing an adiabatic transformation allows the system to be expressed in the qubit subspace (when considering the traditional scenario of leakage into a qubit's third energy level) and, thus, enable easy removal of unwanted leakage. Starting with a Hamiltonian in the lab frame

$$H = \hbar \sum_{j=1,2} [\omega_j \Pi_j + \mathcal{E}(t) \lambda_j (\sigma_j^+ + \sigma_j^-)] \quad (70)$$

by applying the rotating wave approximation (i.e. ignore terms which oscillate very rapidly) [27] and moving into the drive frequency's interaction frame, the

Hamiltonian becomes

$$H^R = \sum_{j=1,2} [\delta_j \Pi_j + \frac{\mathcal{E}^x(t)}{2} \lambda_j \sigma_{j-1,j}^x + \frac{\mathcal{E}^y(t)}{2} \lambda_j \sigma_{j-1,j}^y] \quad (71)$$

where ω_j is the frequency of the j^{th} energy level (ground energy level set to zero). $\Pi_j = |j\rangle\langle j|$ and $\sigma_j^- = |j-1\rangle\langle j|$ and $\sigma_j^+ = |j+1\rangle\langle j|$ are the projector and ladder operators for the j^{th} level, respectively. $\sigma_{j,k}^x = |k\rangle\langle j| + |j\rangle\langle k|$ and $\sigma_{j,k}^y = i|k\rangle\langle j| - i|j\rangle\langle k|$. λ_j denotes the "relative strength of the 0-1 and 1-2 transition" such that $\lambda_1 = 1$ and $\lambda_2 = \lambda$ [8]. δ_j is the difference between the frequency of the j^{th} level and the drive frequency. If the drive and qubit frequencies are resonant, then $\delta_1 = 0$ and $\delta_2 = \omega_2 - 2\omega_1$, which is a feature of the qubit known as the anharmonicity, η [8]. Anharmonicity measures how nonlinear a qubit's energy levels are. The drive and control is described by $\mathcal{E}(t) = \mathcal{E}^x(t)\cos(\omega_d t) + \mathcal{E}^y(t)\sin(\omega_d t)$ when $0 < t < t_g$ where t_g is the time for running one gate. In the situation of a NOT gate, which is of particular interest to this project as well as what was presented in Motzoi, Gambetta, Rebentrost, et al. (2018) [8], $\mathcal{E}^x = \mathcal{E}_\pi$ and $\mathcal{E}^y = 0$, where $\int_0^{t_g} \mathcal{E}_\pi = \pi$. For a Gaussian pulse, this becomes $\mathcal{E}_g(t) = A \exp\left\{-\frac{(t-\frac{t_g}{2})^2}{2\sigma^2}\right\} - B$, where σ is the standard deviation and B ensures the pulse starts and ends at zero [8].

The adiabatic transformation $V(t) = \exp\left\{-i\mathcal{E}^x \frac{\sigma_{0,1}^y + \lambda\sigma_{1,2}^y}{2\eta}\right\}$ allows analysis of the system in the qubit subspace [8]. When using DRAG, the pulses must always start and end at zero, or rather $V(0) = V(t_g) = \mathbb{I}$, which is done by forcing $\mathcal{E}^x(0) = \mathcal{E}^x(t_g) = 0$. The effective Hamiltonian under the adiabatic transformation is now

$$H^V = V H^R V + i\hbar \dot{V} V^\dagger \quad (72)$$

or, fully written out, this can be expressed as follows.

$$\begin{aligned} \frac{H^V}{\hbar} \approx & \frac{\mathcal{E}^x}{2} \sigma_{0,1}^x + \frac{\lambda \mathcal{E}_x^2}{8\eta} \sigma_{0,2}^x + (\delta_2 + \frac{(\lambda^2 + 2)\mathcal{E}_x^2}{4\eta}) \Pi_2 + \\ & (\delta_1 - \frac{(\lambda^2 - 4)\mathcal{E}_x^2}{4\eta}) \Pi_1 + [\frac{\mathcal{E}^y}{2} + \frac{\dot{\mathcal{E}}^x}{2\eta}] (\sigma_{0,1}^y + \lambda \sigma_{1,2}^y) \end{aligned} \quad (73)$$

The last two terms in equation 73 reflect the leakage, thus by setting $\delta_1 = \frac{(\lambda^2 - 4)\mathcal{E}_x^2}{4\eta}$ and $\mathcal{E}^y = -\frac{\dot{\mathcal{E}}^x}{\eta}$, the leakage can be negated. Of particular note is the last term, which explicitly shows how information can transfer into $|2\rangle$ by the presence of $\sigma_{1,2}^y$. In this way, by adding a second control, $\mathcal{E}^y$, proportional to the derivative of the original pulse, leakage can be "instantaneously" reduced [8]. While a Gaussian shape was used in this project, if future work determines there is a better pulse shape, DRAG may still be useful as it "is largely impervious to envelope shape, provided it has the right area and starts and ends at 0" [8].

5 Implementation

5.1 Mean Field Approximation

This project simulated a two-dimensional lattice of LC resonators in an arrangement that mimics the pendula studied in Süsstrunk and Huber (2015) [6]. The resonators, much like the pendula, are dictated in pairs, or unit cells. The lattice is 45 by 45 unit cells or 90 by 45 resonators. If this lattice were to be fabricated and tested, each individual coupling between every resonator would need to be designed. This is an inefficient method for simulation purposes, so instead the mean-field approximation is introduced in the numerical analysis [28]. This allows us to analyze the system in simulation without tuning up every individual coupling.

The system Hamiltonian for a driven lattice with one qubit attached to one of the resonators is generalized to

$$H = \sum_i \hbar \omega_r a_i^\dagger a_i + \hbar \frac{\omega_q}{2} \sigma_z + g_{r',q} (a' \sigma_+ + a'^\dagger \sigma_-) + f(t) (a_d + a_d^\dagger) \quad (74)$$

where ω_r is the resonator frequency, a_i and $a_i^\dagger$ are the raising and lowering operators of resonator i with a' and $a'^\dagger$ denoting the raising and lowering operators for the specific resonator r' , which is connected to the qubit. ω_q is the frequency of a qubit connected to resonator r' , $g_{r',q}$ is the coupling between the resonator and the qubit, $f(t)$ describes the drive and a_d and $a_d^\dagger$ are the raising and lowering operators for the resonator that is initially driven by the pulse. The mean-field approximation dictates that the time derivative of any of these ladder operators for resonators or qubits is defined as

$$\frac{da_i}{dt} = \frac{i}{\hbar} [H, a_i] \quad (75)$$

This also applies to the qubit operators, such that

$$\frac{d\sigma_i}{dt} = \frac{i}{\hbar} [H, \sigma_i] \quad (76)$$

where $i \in \{z, +, -\}$. Note that in the numerical analysis conducted, $\hbar = 1$. By stepping through each operator for each component in the Hamiltonian, the system's time derivative can be integrated over. This is done using the packaged `odeintw`, a wrapper around SciPy's `integrate.odeint`. `Odeint` finds the integral of ordinary differential equations. `Odeintw` enables the same calculations on complex-valued systems [29]. While the transformation from $\mathcal{H}$ to $\mathcal{D}$ introduced by Süsstrunk and Huber (2015) puts the complex Hofstadter's model into a real space [6], qubits by nature are complex and so including even a single qubit in the Hamiltonian makes the system at large complex, again. For this reason, SciPy's `odeint` is insufficient as it does not handle complex differential equations well. The code for this project can be found at <https://gitlab.tudelft.nl/qmai/master-theses/genyacrossman>.

5.2 The Lattice

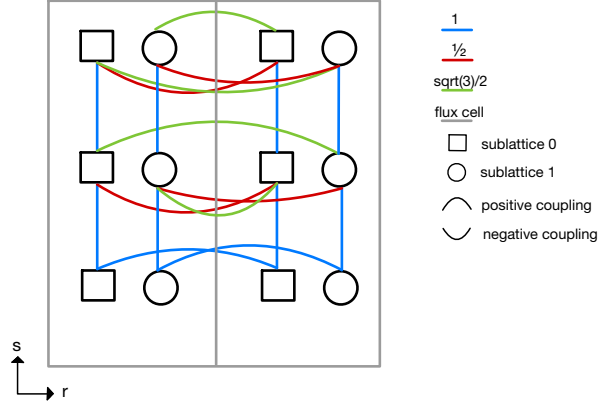


Figure 15: The couplings among resonators within two flux cells. Each square (sublattice 0, i.e. pendulum x in [6]) and circle (sublattice 1, i.e. pendulum y in [6]) form a unit cell. All squares form sublattice 0 and all circles form sublattice 1. Three unit cells together in the s direction form a flux unit cell. In other words, the flux among three unit cells is 2π . This is a property of the Hofstadter model. All couplings are proportional to $f = \frac{1}{2\pi} \sqrt{\frac{1}{LC}}$, as was introduced in equation 14. The blue lines show couplings equal to f , red lines represent couplings that are one half of f , and green lines are $\frac{\sqrt{3}}{2}f$, which is the only coupling connecting the two sublattices. Upward bending arcs represent positive couplings, whereas downward arcs are negative couplings. All blue couplings are positive.

In an effort to mimic the pendula system presented in Süsstrunk and Huber 2015 [6], the couplings between the resonators in the array are implemented in exactly the same way, as shown in figure 15. In the pendula system, f is related to the moment of inertia and torque constant of the springs [6]. In this circuit scenario, f is related to the inductance and capacitance of the resonators. The couplings between two flux cells – i.e three unit cells or pairs of resonators – are described as follows: the coupling between resonators of the same sublattice in a single flux cell are all equal to f (see blue lines in figure 15). This is also the coupling connecting the resonators of the same sublattice in two neighboring flux cells at the lowest r . The couplings between resonators of the same sublattice in neighboring flux cells of higher r values is proportional to $\frac{-f}{2}$. Figure 15 denotes these couplings in red. The downward arc of the connections represents the fact that these couplings are proportional to negative f . Lastly, the only way the two sublattices are connected is through $\frac{\sqrt{3}f}{2}$ – both positive and negative, again, reflected in the upward or downward arc of the connection, respectively. These

connections are also only on the two higher r unit cells (see the green lines in figure 15).

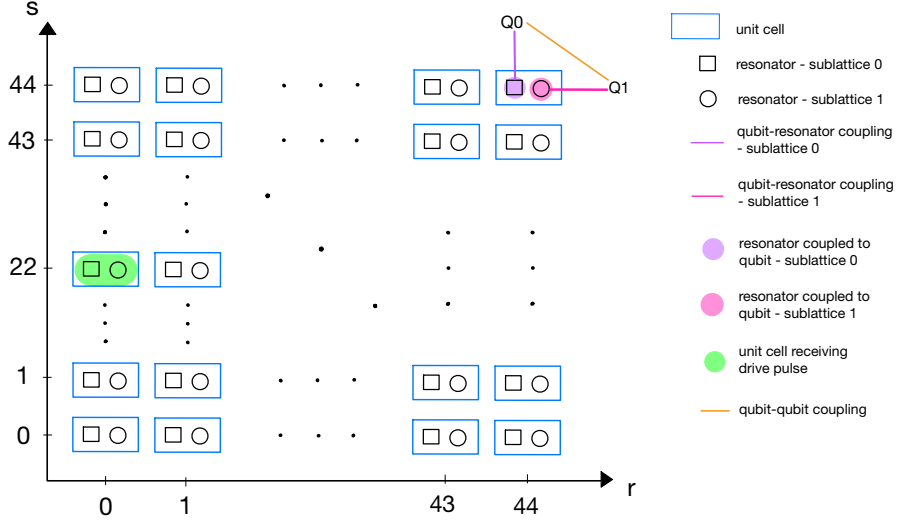


Figure 16: The full LC resonator lattice with two qubits attached. In the one-qubit study, only qubit 1 is used; qubit 0 is not introduced until the two-qubit system. The pulse enters at lattice site (0,22).

The lattice studied is comprised of 45 unit cells, or pairs of resonators, which means 90×45 resonators, as in the r direction there are two resonators for every unit cell, but only one in the s direction, as shown in figure 16. Thus, each column in $\hat{s}$ is comprised of resonators all from the same sublattice, whereas every row in $\hat{r}$ contains elements from both sublattices. Moving forward, an individual resonator will be denoted as (r, s, l) where r and s are its locations in $\hat{r}$ and $\hat{s}$, respectively, and $l \in \{0, 1\}$ denotes which sublattice the resonator belongs to. If only two indices are given, e.g. (3,25), then it implies a unit cell (both resonators) at this location. The orientation has been chosen such that the left-hand resonator in a pair is from sublattice 0 and the right-hand from sublattice 1. This means that the start of the lattice at the bottom left is a sublattice 0 resonator, (0,0,0), and the top right of the lattice is a sublattice 1 resonator, (44,44,1). When a single qubit is added, it is connected to the top-right most resonator (44,44,1). When considering a second qubit, it is attached to (44,44,0) so that it can easily interact with the original qubit from sublattice 1. The system is excited with a pulse at (0,22). Exactly how the system is excited will be discussed next.

5.3 Pulse Shaping

Previously, a Gaussian pulse was loosely tuned up for this lattice size such that the initial excitement from the pulse finished prior to the subsequent wave hitting a corner of the lattice. Such Gaussian parameters are t_0 of 10ns, $2\sigma^2$ of 36, and an amplitude, A , of 1. As the immediate goal of this work is not to fine tune the pulse parameters, but to understand at a higher level the interaction of the pulse and lattice, changing these parameters were not within scope of this work. Thus, the original Gaussian pulse, $f_d(t)$, can be described as it's effect on each sublattice, f_d^0 and f_d^1 , as

$$f_d^0(t) = A \exp\left\{-\frac{(t-t_0)^2}{2\sigma^2}\right\} \sin\left(\frac{\theta}{2}\right) \sin(\omega_d t) + \cos\left(\frac{\theta}{2}\right) \sin(-\omega_d t + \varphi) \quad (77)$$

$$f_d^1(t) = A \exp\left\{-\frac{(t-t_0)^2}{2\sigma^2}\right\} \sin\left(\frac{\theta}{2}\right) \cos(\omega_d t) + \cos\left(\frac{\theta}{2}\right) \cos(-\omega_d t + \varphi) \quad (78)$$

where ω_d is the frequency of the pulse and θ and φ are the polarization angles. θ works such that when it is zero, the signal travels clockwise, when it's π the signal travels counterclockwise, and when it is $\frac{\pi}{2}$ it is split and goes in both directions. θ was kept at 0 in these simulations. This means that to enact an X gate, for example, the signal effectively needs to flip it's θ by π . Or, if an H gate is desired, then the signal needs to be changed somehow by $\theta = \frac{\pi}{2}$. Understanding how to do this requires a sense of how to measure when the edge mode flips in chirality. Such metrics are discussed in the following section.

Prior to exploring this, though, when looking at the whole lattice with a qubit, an unwanted backpropagation of the signal is seen after the qubit interacts with the excited lattice. Thus, prior to considering any methods for enacting a gate, first understanding how to get a strong initial continuous signal is important. To do this, a DRAG inspired pulse was used (see chapter 4). GRAPE could be implemented, but it does not appear to require the pulse start and stop at 0. This system requires a temporary initial excitation to the system, not a continuous one. Additionally, the result of GRAPE depends on the discretization of the signal in time. As there are already many unknowns about how parameters interact in this system, introducing another uncertainty at this stage of developing a toy-model could threaten increasing the complexity of the case. DRAG appears to be the ideal candidate for this scenario given its design for minimizing information leakage into unwanted states, but it heavily relies on the qubit-specific parameter anharmonicity, which does not exist for linear elements like LC resonators. DRAG does, however, remove the concern of discretization presented by GRAPE and ensures the signal starts and stops at zero through its "instantaneous" addition of terms [8]. As the point is also not to fully optimize the system as is exactly presented currently, but merely demonstrate that backpropagation can be minimized, a simple method of adding the derivative of the Gaussian pulse (introduced in equations 77 and 78) to the Gaussian itself was implemented. So long as the backpropagation is minimized, then the rest of the system can be more thoroughly studied. Once the system is

better understood (e.g. the architecture and Hamiltonian parameters are well defined), more thorough optimal control methods, such as GRAPE or DRAG, should be implemented as is necessary. Going forward, the original Gaussian pulse will be referred to as the "unoptimized" pulse whereas the Gaussian with its derivative will be called the "optimized" case.

To understand to what extent this method changes the system, several perspectives were considered. Of course, first off, the backpropagation was studied to see if its effect was minimized. Then how the signal strength varied from the edge of the lattice in comparison to the center was considered; were edge modes maintained even with an optimized pulse? How does the signal after the qubit compare to the signal before the qubit? And how does adding the pulse's derivative impact a system when 2 qubits are involved? The following chapters will present the results to address these questions.

5.4 Metrics

When comparing quantum states, the metric to determine how much information was preserved in transitioning from one state to the next is referred to as fidelity. For pure states, this calculation is simply $F = |\int \langle \Psi_{start} | \Psi_{end} \rangle|^2$ where Ψ_{start} and Ψ_{end} refer to the starting and ending quantum states. While the edge modes of the lattice are expected to mimic quantum states, the fidelity calculation is not an optimal choice for determining how much of the signal passes the qubit and continues onward as the lattice is not producing pure states. Calculating the fidelity for mixed (non-pure) states requires knowledge of the system that may be commonly accessible in traditional quantum systems, but is not currently in this setup. There is, however, another metric common in quantum systems, as well as in microwave engineering.

The transmission coefficient is often introduced in introductory quantum mechanics as a method for determining how much of a wave transmits past a potential barrier [25]. In microwave engineering, it is similarly used to determine how much a wave transmits past a specific circuit element, such as a resonator [30]. In essence, it is a description of how much a wave continues past some barrier. For simplicity, consider simple plane waves, such that $\Psi_L = Ae^{ikx} + Be^{-ikx}$ is the expression of the wave to the left of the barrier and $\Psi_R = Ce^{ikx} + De^{-ikx}$ is the wave to the right of the barrier. Both Ψ_L and Ψ_R contain information about the ingoing and outgoing waves. The system can be rewritten as $\Psi_{in} = Ae^{ikx} + De^{-ikx}$ and $\Psi_{out} = Ce^{ikx} + Be^{-ikx}$. Then consider the situation in which the wave approaches the barrier from the left-hand side, therefore there is no incident wave approaching the barrier from the right and, thus, $D = 0$. As the transmission coefficient is simply a comparison of how much of the original incident wave continues onwards after the barrier, it can be expressed in terms of the amplitudes, A and C [25].

$$TC = \frac{|C|^2}{|A|^2} \quad (79)$$

This means that so long as the amplitudes of the wave before and after the

barrier are known, one can determine how much of the wave was transmitted. This works wonderfully for the lattice system, if the qubit is considered a barrier. Understanding how much of the original signal is transmitted past the qubit will help determine how robust the system is. For example, if an X gate was run, ideally there should be none of the original wave passing the qubit; it should all be reflected. The transmission coefficient in this case ideally would be zero. In this situation, the reflection coefficient, $RC = \frac{|B|^2}{|A|^2}$, would be 1 as the wave is fully reflected and $TC + RC = 1$ [25]. If, however, the qubit was not supposed to interfere with the system, then the ideal transmission coefficient would be 1. While the transmission coefficient itself is a very good metric for this system and its future incarnations, determining exactly which amplitudes to use is not immediately obvious. A discussion around this can be found in chapter 6.

6 One Qubit

The first step to checking how the system behaves is analyze the behavior of the edge modes. Verifying that the signal decays as it progresses away from the edges and into the bulk is essential before considering a more complex system with multiple qubits. Thus, first results presented relate to characterizing the edge signal versus the bulk signal. This includes studying how an optimized pulse may impact the system. Then the a single qubit's response is considered in both the unoptimized and optimized scenarios for when the qubit is attached at the top right corner of the lattice on sublattice 1, i.e. at resonator (44,44,1). The qubit frequency is set to the drive frequency; both are 2.2GHz. These two features were tested in resonance as that is the point at which the qubit is likely to be maximally excited. It is unknown, yet, whether an LC lattice fully functioning as a qubit should have its control qubit maximally or minimally excited. This will likely depend on the type of gate being run. By studying the maximally excited control qubit, however, the extreme condition of qubit excitation is better understood. Section 6.3 covers the results of the transmission coefficient – as was introduced in section 5.4 – for both the unoptimized and optimized pulses. These are shown for time scales of 1σ , 3σ , and 5σ centered about the respective peaks, as well as for 0 to 100ns. The transmission coefficient is calculated twice: first comparing four unit cells before and after the qubit then comparing the center of the top edge of the lattice to the center of the right edge of the lattice. A discussion around the limitations of these results is also presented. Finally, the qubit-resonator coupling was studied to ensure that this parameter is not limiting the Hamiltonian prior to complicating the system by adding a second qubit. Indeed, the qubit-resonator coupling, g , is not currently a limiting factor in the current setup, so it is considerably safe to add a second qubit to the lattice. Two-qubit results are presented and discussed in chapter 7. For the following results, the drive's amplitude is 1.0, its phases, θ and φ , are both zero, $t_0 = 10\text{ns}$, and $\sigma = \sqrt{18}$. The qubit-resonator coupling is 0.8 GHz.

6.1 Edge Modes

As mentioned, adding a qubit to the lattice – attached to sublattice 1 at the top right of lattice, i.e. resonator (44,44,1) – most notably caused a backpropagation along the edge of the lattice. This was mitigated using an optimized pulse, as was introduced in section 5.3), and which will be discussed at the end of this section. Adding a qubit, however, also served as a check to see how the topological edge modes hold up when the edge shape changes. In Süsstrunk and Huber (2015), this was tested by removing pendula [6]. With integrated circuits, though, LC resonators are well defined and easily fabricable components, which is a large part of the motivation in using them for this research. Removing resonators, thus, would not mimic a truthful experimental change in the definition of an edge, as resonators are not likely to "disappear" from a device. What is more likely, is that the edge is extended by being attached to another element, such as a qubit. Attaching a single qubit to the system

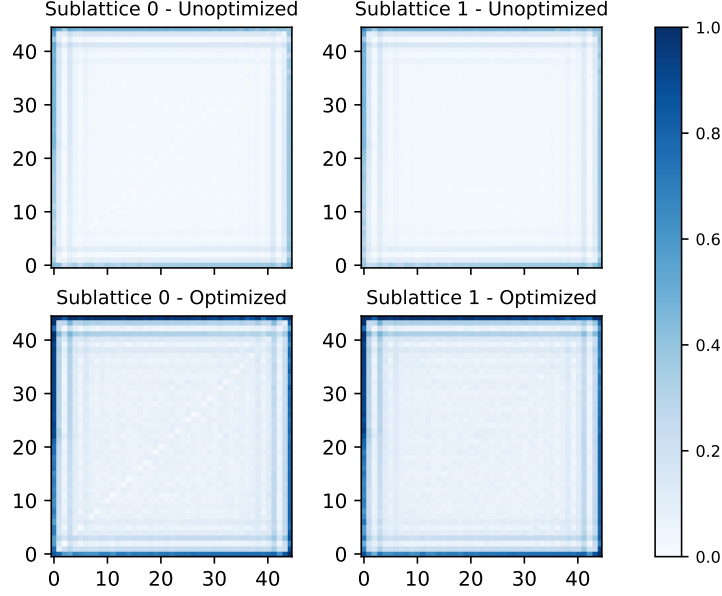


Figure 17: The average amplitude of each resonator of a given sublattice over all time when one qubit is attached to sublattice 1 at the top right corner, i.e. (44,44,1). Results are shown for both sublattices under two different conditions: an unoptimized pulse and an optimized pulse.

physically extends the edge of where the signal can travel, while also tests out the foundations of the lattice-with-qubit dynamics in anticipation of a two-qubit system. As expected, attaching the qubit does not disturb the observation of topological edge modes (see figures 17 and 18²). There is a very slight diagonal signal that travels through the bulk from where the qubit sits in the top right corner to the bottom left corner of the lattice after the qubit becomes excited. This is present with or without the pulse's derivative, but is more visible with it as the optimized pulse amplifies the entire signal. Nonetheless, the strength of this diagonal wave is negligible in comparison to the main signal, so it can be ignored for the purposes of this project. Also, as the qubit is set at the same frequency of the drive, it is maximally excited; therefore, future studies of this system with different qubit frequencies will likely demonstrate a lower diagonal bulk signal, if not remove it entirely.

²Gif files for these results are in the Images folder under the file names *Unopt_sublattice0_Qat(44,44,1)*, *Unopt_sublattice1_Qat(44,44,1)*, *Optimized_sublattice0_Qat(44,44,1)*, *Optimized_sublattice1_Qat(44,44,1)*, *Unopt_UnitCell_Qat(44,44,1)*, and *Optimized_UnitCell_Qat(44,44,1)*. These can be found at <https://gitlab.tudelft.nl/qmai/master-theses/genyacrossman>

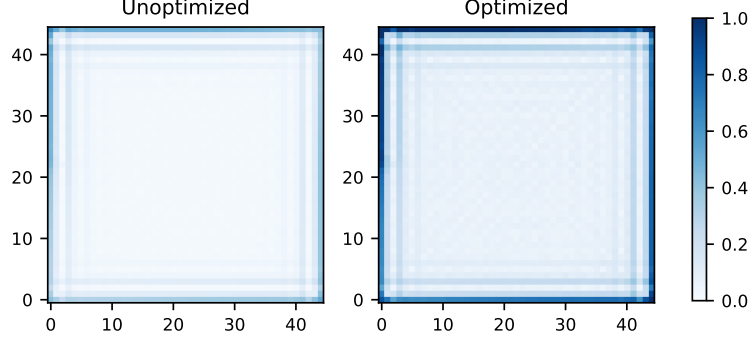


Figure 18: The average amplitude of each unit cell over all time when one qubit is attached to sublattice 1 at the top right corner, i.e. at $(44,44,1)$. One unit cell's amplitude is the average of the responses of the two resonators – one from sublattice 0 and one from sublattice 1 – that comprise that cell. Results are shown for under two different conditions: an unoptimized pulse and an optimized pulse.

What matters more is how the signal decays into the bulk or how it behaves along the edge after it hits the qubit. The signal drops by about one half when only one resonator in from the edge, regardless of sublattice. This decay is not linear, however, as can be seen in figures 17 and 18. Roughly every 3 layers in from the edge, there is a slight increase in the signal on an order of magnitude similar to about 3 layers previous. While the particular dynamics of this nonlinear decay into the bulk have not yet been characterized, it is not concerning as the magnitude of the wave amplitude anywhere in the bulk is still significantly less than (at least half of) that at the edge.

While the edge modes exhibit similar robustness given both the unoptimized and optimized pulse conditions, the effect of the backpropagation along the top edge of the lattice caused by the qubit differs between these two conditions. Figure 19 shows the response of resonator half-way along the top edge and on the same sublattice as the qubit, i.e. resonator $(22,44,1)$. The dashed yellow and orange vertical lines denote when the qubit's amplitude is greater than 10^{-10} and 10^{-5} respectively. Thus, as expected, this resonator is maximally excited before the qubit's response is substantial. This initial peak is part of what looks like the only clean Gaussian in all time for both the unoptimized and optimized scenarios. The rest of the response is comprised of the backpropagation from the qubit then the subsequent moments the main pulse hits same resonator again after making its way around the lattice and the related backpropagations from each of those rotations. Considering the signal gets messy at later points

in time and that these later time frames are likely irrelevant to this project³, a resonator's response will only be considered for the first 100ns, as seen in figure 20. Almost immediately after the initial pulse passes the resonator, right when one would hope the resonator would then become relatively dormant, there is another messier signal that passes through. This is easier to see in figure 20b, which shows the same resonator response, but only for the first circulation of the signal around the lattice. The optimized pulse (light blue line in the figures) does not get rid of the backpropagation in this system, rather it diminishes its continuity of amplitude in comparison to the unoptimized pulse. More importantly, the backpropagation amplitude is significantly less in comparison to the peak of the main signal from an optimized pulse. In the unoptimized case, the backpropagation is nearly the same amplitude as the original signal. In the optimized case, the backpropagation is not even half of the maximum amplitude of the original signal. So while the optimized pulse does not make backpropagation disappear for the lattice with a maximally excited control qubit, it does make it negligible in comparison to an unoptimized pulse. Additionally, the backpropagation may diminish in future studies if the qubit and drive frequencies need to be detuned.

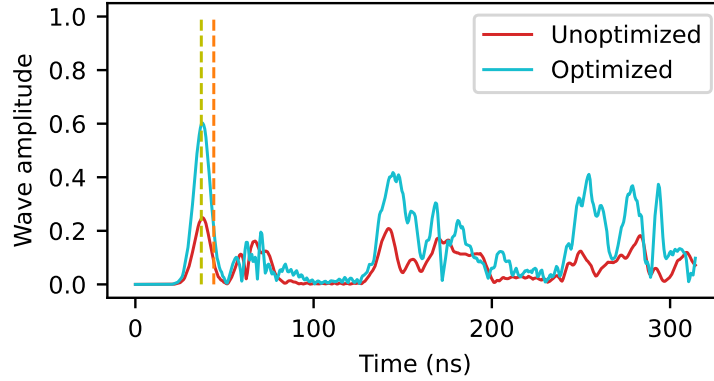
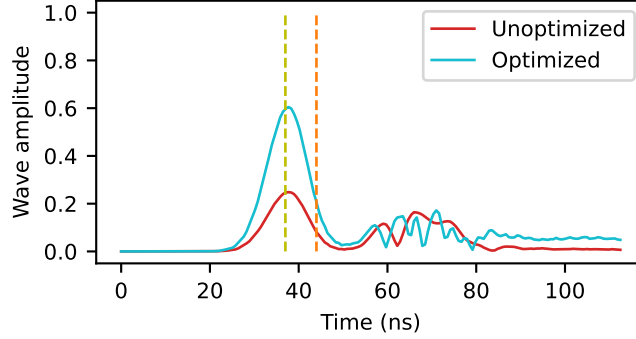
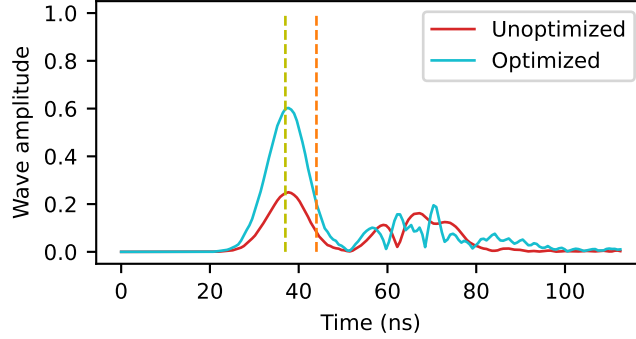


Figure 19: Sublattice 1 response at half-way across the top edge, i.e. at (22,44,1), over all time. The yellow and orange dashed vertical lines show when the qubit is excited to an amplitude greater than 10^{-10} and 10^{-5} , respectively.

³Current two-qubit gates in superconducting systems are on the order of magnitude of 100ns [26]



(a) Sublattice 0



(b) Sublattice 1

Figure 20: Responses of the two resonators – one of each sublattice – at the center of the top edge. a) Resonator of sublattice 0, i.e. $(22,44,0)$. b) Resonator of sublattice 1, i.e. $(22,44,1)$. Yellow and orange dashed vertical lines denote when the qubit’s amplitude is at least 10^{-10} and 10^{-5} , respectively. The original pulse’s signal is clearly seen as the initial Gaussian and the remainder of what looks like noise is the backpropagation from the qubit.

6.2 Qubit Response

After checking the edge modes, it was important to ensure the qubit is being properly excited. While this project did not get into the specifics of optimizing the qubit behavior, it is still crucial to understand that the qubit is “alive,” or rather that it is observed oscillating between the ground and the excited state. Figure 21 shows the qubit attached to the lattice that is excited by an unoptimized pulse. The qubit starts in the ground state. Once it is significantly excited, shortly before 50ns, it begins to oscillate. It does not, however, return fully to the ground state; rather, it appears to oscillate between the almost-excited state and a mixed state, i.e. some state between the ground and excited state. Similarly to the resonator’s argument of only considering the first 100ns of time, here it is more obvious how at longer time scales the quality of the control

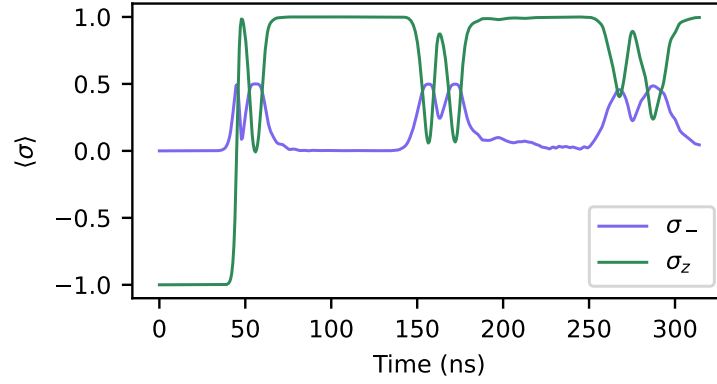


Figure 21: The control qubit's response for a system excited by an unoptimized drive.

qubit decays as its oscillations become less pronounced. Figure 22 shows the qubit responses to both unoptimized and optimized pulses on this time scale. While it was not expected for the optimized pulse to impede any qubit activity, it is good to confirm this (see figure 22b). It also appears that the optimized pulse helps the qubit get closer to oscillating fully back to its ground state, which is what would ideally happen after running a gate – the control qubit would return to its initial starting point. Understanding the dynamics of why the control qubit has one extra oscillation with an optimized pulse requires further investigation.

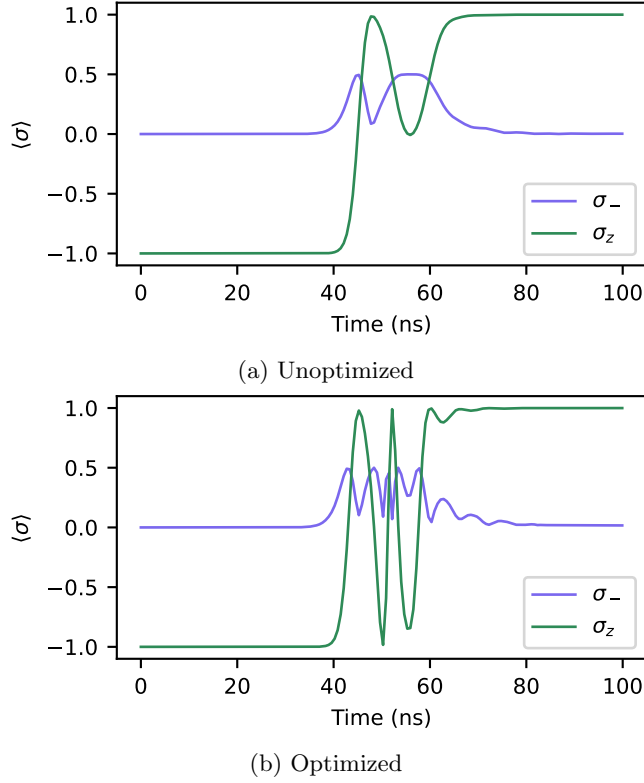
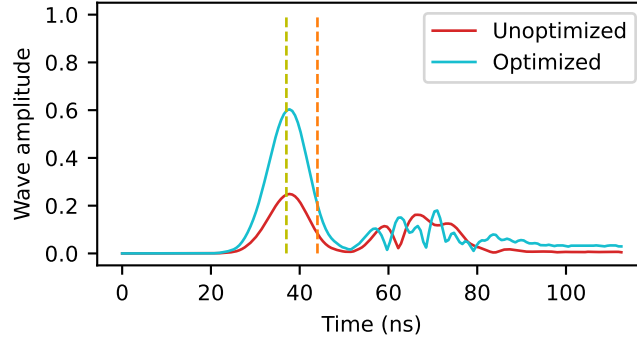


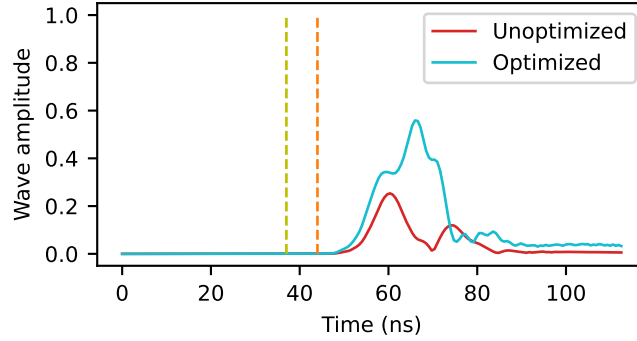
Figure 22: The control qubit's response for a system excited by a) an unoptimized or b) optimized pulse.

6.3 Transmission Coefficient

As introduced in section 5.4, the transmission coefficient is a good metric for determining how much of a wave has continued past a barrier. In this case, the barrier is the control qubit. While taking the maximum amplitude of a given unit cell before the qubit and comparing it with the maximum amplitude of a given unit cell after the qubit may seem sufficient enough for calculating the transmission coefficient of the system, it does not take into account how the shape of the signal changes after the qubit. This change in shape is particularly pronounced when using an optimized pulse, see figures 23 and 24. For this reason, the integral of the amplitude of a specific unit cell's response was taken and used instead of a single amplitude at a single point in time. It is not obvious exactly what time frame to integrate over, however. Particularly for any sites on the top edge, it will be difficult to find a window of time at which it is guaranteed there is no backpropagation being accidentally added into the integral of the "original" wave. Thus, the integral was calculated four times, each for different time frames. At first the time-width considered equated to σ , 3σ , and 5σ , where



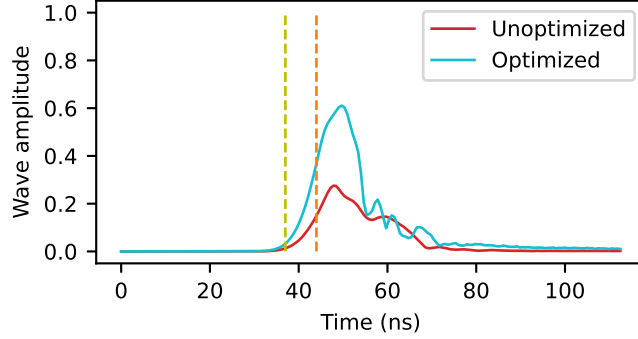
(a) Half-way across top edge



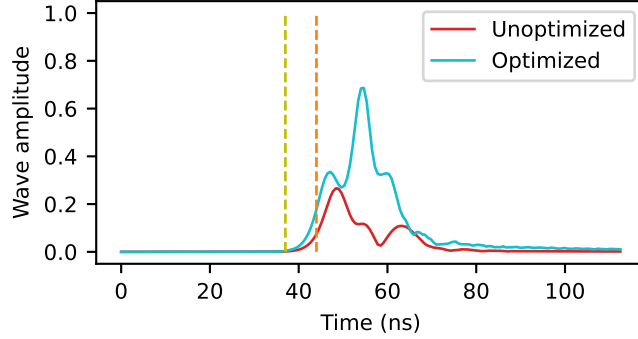
(b) Half-way across right edge

Figure 23: Unit cell responses at the center of the lattice before and after the qubit, i.e. at (22,44) and (44,22), respectively. Yellow and orange dashed vertical lines denote when the qubit's amplitude is at least 10^{-10} and 10^{-5} , respectively.

σ refers to the standard deviation of the Gaussian of the drive, centered about the respective peaks of the lattice sites. Please note, though the optimized pulse incorporates the derivative of the Gaussian (so its standard deviation may not be exactly σ), these same time frames were still used for the sake of simplicity and getting ballpark estimates. Luckily, the transmission coefficient is independent of the standard deviation, this is merely a practice for appropriately highlighting the part of the signal of interest, therefore even if the optimized pulse has a different standard deviation, this should not wildly affect the results. The fourth time frame considered was 0 to 100ns. These four time frames were used to integrate the signal at four different lattice sites to calculate two transmission coefficients for each combination of time frame and pulse type (unoptimized and optimized). Table 1 shows the resulting transmission coefficients for comparing four lattice sites before the qubit, (40,44), with four lattice sites after the qubit, (44,40), corresponding to figure 24. Table 2 shows the transmission coefficients



(a) Half-way across top edge, i.e. before qubit.



(b) Half-way across right edge, i.e. after qubit.

Figure 24: Unit cell responses at four lattice sites before and after the qubit, i.e. at (40,44) and (44,40), respectively. Yellow and orange dashed vertical lines denote when the qubit's amplitude is at least 10^{-10} and 10^{-5} , respectively.

for comparing lattice sites at the center of the edge before the qubit, (22,44), and the center of the edge after the qubit, (44,22), corresponding to figure 23.

First, it is important to note that as expected, the unoptimized transmission coefficients are less than the optimized results. The exception to this is with the 1σ values when comparing lattice sites at the center of the edges before and after the qubit (see table 2). The fact that the unoptimized transmission coefficient is larger here indicates that 1σ is too small of a window. The rest of the numbers, however, generally agree with the idea that the optimized pulse enables more of the signal to continue onwards past the control qubit. The 0 to 100ns results for both four unit cells from the qubit and centers of the edges show roughly the order of magnitude expected. This time frame, however, is not perfect. It likely could contain the backpropagation into the "before" signal, especially in the study of four unit cells away from the qubit as presented in table 1. Ideally, the backpropagation should be extracted from any "before" amplitudes prior to any further calculations for the purposes of the transmission coefficient. This

	Time slice			
	1σ	3σ	5σ	0 to 100 ns
Unopt.	0.89	0.6	0.5	0.52
Opt.	1.07	0.82	0.98	0.93

Table 1: Transmission coefficients for both unoptimized and optimized solutions four unit cells before and after the qubit, i.e. (40,44) and (44,40), for varying lengths of time. σ , 3σ , and 5σ are centered about the respective peaks.

	Time slice			
	1σ	3σ	5σ	0 to 100 ns
Unopt.	0.98	0.75	0.69	0.35
Opt.	0.8	0.86	1.07	0.74

Table 2: Transmission coefficients for both unoptimized and optimized solutions half-way across the top and right edges, i.e. (22,44) and (44,22), for varying lengths of time. σ , 3σ , and 5σ are centered about the respective peaks.

would make the results more accurate. Hopefully such improved transmission coefficients would be no more than 1, as the current values greater than 1 are not physical. As mentioned in section 5.4, a transmission coefficient of 1 implies that the entire wave was transmitted past the qubit. A transmission coefficient larger than 1 here likely indicates that both the time frame considered is not the best and/or that the backpropagation is erroneously included in the "before" wave amplitude. It also possibly indicates that a further normalization step is likely required after the integration of the signal and before the final transmission coefficient computation.

6.4 Qubit-Resonator Coupling

Lastly, the qubit-resonator coupling was varied to ensure that this parameter is not limited while in the simplest case of connecting a single qubit to the lattice. At first, the unoptimized system was varied from $g = 0.2$ to $g = 1.2$ GHz. After seeing that the qubit was too weakly coupled to the system to properly oscillate between the ground and excited state when $g < 0.6$ GHz, the optimized system was considered for $g = 0.6$ to $g = 1.2$ GHz. Both the qubit and the resonator attached to the qubit were considered for each coupling strength to ensure their respective responses were not negatively impacted. After confirming that the Hamiltonian for a lattice with one control qubit is not limited with respect to the qubit-resonator coupling strength, a two-qubit setup was studied. It is important to note that while the qubit-resonator coupling strength was studied, it was not optimized. Full-system Hamiltonian optimization is a complex process; tuning up one parameter, such as the qubit-resonator coupling, in isolation does not optimize the system as a whole. For example, adding a second qubit will further constrain the Hamiltonian and, thus, it is not useful to optimize

the qubit-resonator coupling in a single qubit system when the anticipated final system will contain two qubits.

7 Two Qubits

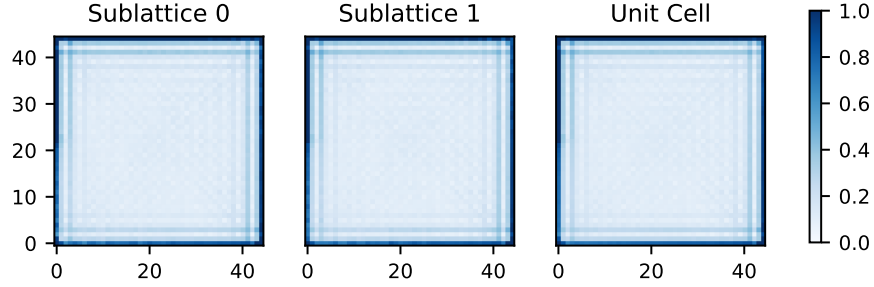
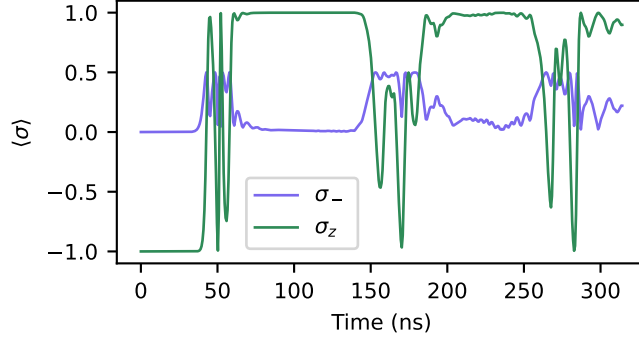


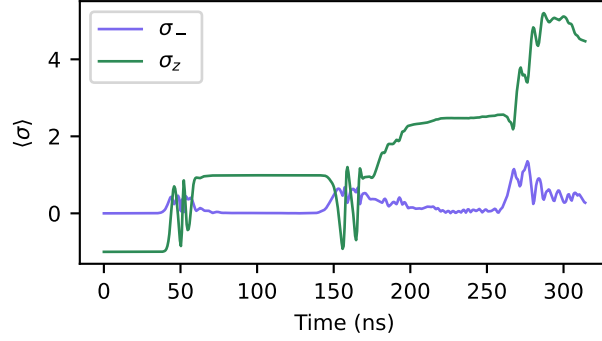
Figure 25: Average signal across the lattice of resonators when two qubits are attached at the top right corner and the system is excited by an optimized pulse. One qubit is attached to sublattice 0 (at resonator $(44,44,0)$) and the other is attached to sublattice 1 (at resonator $(44,44,1)$). The unit cell result shows the average of the two sublattices' responses.

Once understanding that a single qubit can be attached to a lattice of LC resonators without wildly disturbing the edge modes, a system with two qubits was studied. Two qubits will likely be necessary for enacting gates on the lattice as both sublattices need to be controlled. It is logical, then, to consider one qubit attached to each sublattice and the two qubits coupled together. In this section the results of such a system are presented with qubit 0 – the qubit attached to sublattice 0 – having a frequency of 2.2 GHz and qubit 1 – the qubit attached to sublattice 1 – having a frequency of 2.0 GHz. All other parameters remain the same as in chapter 6. Only the optimized pulse was studied as the one-qubit results indicate that this minimizes the affects of backpropagation and helps the control qubits to oscillate closer to the ground state.

In a similar method to exploring a one-qubit system, first the edge modes of the lattice were considered. Again, the signal was found to decay in a non-linear manner into the bulk, but significantly enough such that any bulk-signal is negligible in comparison to an edge-signal (see figure 25). Next, again, much like in the one-qubit case, the excitation of the qubits were considered. As was seen in the one-qubit scenario, the optimized pulse appears to assist the qubits in getting closer to oscillating properly between the ground and excited state. Also, similarly, at longer time scales past 100ns, the quality of qubit oscillation decays for both qubits, but particularly so for the qubit attached to sublattice 1, which has the lower frequency. While the apparent extreme excitation of qubit 1 at longer time scales appears alarming, these are not time scales of concern for



(a) Qubit 0

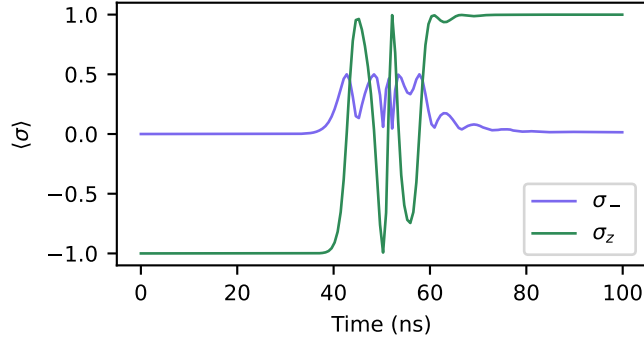


(b) Qubit 1

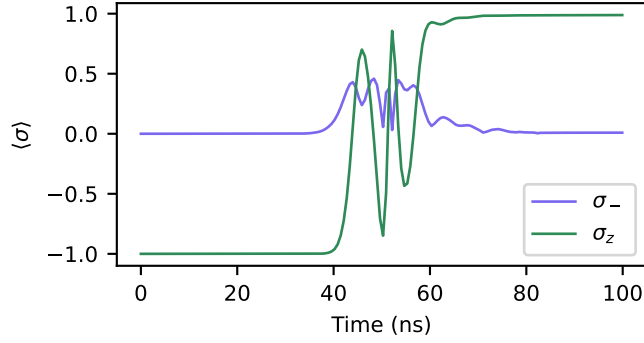
Figure 26: Qubit responses when qubit 0 is attached to sublattice 0 and qubit 1 to sublattice 1. The two qubits are also coupled to each other. Both qubits start in the ground state.

this project. Nonetheless, it should be noted that this behavior may also further highlight the fact that the Hamiltonian has not been fully optimized. Simply applying the optimized pulse does not solve all woes. For example, qubit 1 does not seem to entirely reach its excited state during its first oscillation (see figure 27). Future work must consider the proper allocation in frequency space of the drive pulse, the qubits, and the various couplings.

Figure 28 shows the unit cell response at the center of the top and right edges for a two-qubit system. The shape of these signals and their amplitudes are quite similar to the single qubit case as shown in figure 23. These results serve as good sanity checks that the anticipated eventual two-qubit system will not behave in a wildly different manner than its simpler predecessor, the one-qubit system. With that said, there are still aspects to further study in the single qubit system prior to optimizing for two qubits. For example, given that the proper implementation of the transmission coefficient requires more thorough consideration, it was not implemented here for the two-qubit case



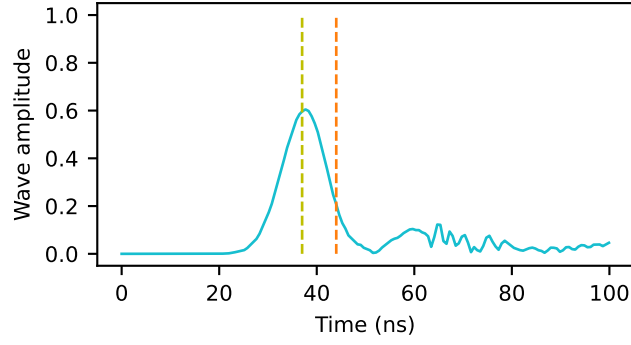
(a) Qubit 0



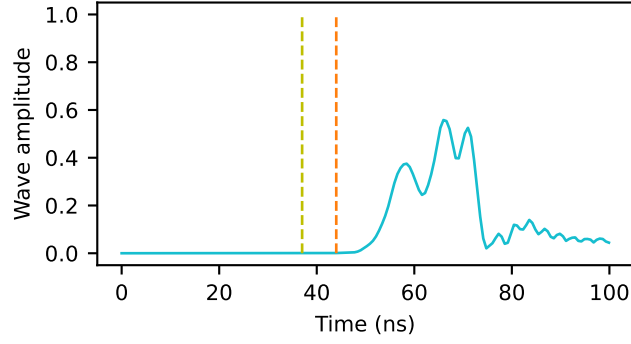
(b) Qubit 1

Figure 27: Qubit responses over the first 100ns when qubit 0 is attached to sublattice 0 and qubit 1 to sublattice 1. The two qubits are also coupled to each other. Both qubits start in the ground state.

even though the transmission coefficient would be a good metric to characterize this system. Such further considerations and suggested next steps are outlined in the following chapter.



(a) Half-way across top edge, i.e. before qubits.



(b) Half-way across right edge, i.e. after qubits.

Figure 28: Unit cell responses at the center of the top, (22,44), and right-hand, (44,22), edges. Yellow and orange dashed vertical lines denote when qubit 1's amplitude is at least 10^{-10} and 10^{-5} , respectively.

8 Conclusion

Studying the proposed LC resonator two-dimensional lattice with control qubit(s) from the perspective of architecture, optimal control, and then metrics produced encouraging results indicating that such a classical array can demonstrate topological-inspired edge modes. The two sublattices were verified to react to an excitation pulse in a similar manner even if only one sublattice was connected to a qubit. When a single qubit was connected to one sublattice, the qubit's state oscillated as one would hope, although does not fully oscillate between the ground and excited state. Additionally, it ends in the excited state, whereas ideally, once a single qubit gate is run, it should end in the ground state where it originally started. This, however, does not necessarily indicate a problem with the layout, moreover it points to a need to find a better relationship between the qubit and the rest of the system; in this work the qubit is in resonance with the drive frequency to become maximally excited such that the system studied is in some type of known "boundary condition" Hamiltonian-wise. Additionally, if the qubit is moved off resonance from the drive pulse, the backpropagation observed in the unoptimized system may diminish. Nonetheless, the backpropagation in this maximally-excited-control-qubit system was minimized by adding the derivative of the Gaussian excitation to the pulse.

While putting the qubit and drive pulse in resonance implies a maximally excited qubit, the qubit-resonator coupling does not necessarily have an prior known maximum (or minimum); however, by changing this variable, it was found that the response of the circuit is not immediately limited by qubit-resonator coupling. This is as one would hope, but was a good check prior to making the system more complex with a second qubit.

When a second control qubit was added to the other sublattice and coupled to the original qubit, the system behaved similarly to the single-qubit system. The oscillations of the lower frequency qubit, however, were further from full ground-to-excited state oscillations. Once again, this indicated the need for optimizing the full Hamiltonian. Optimization of the Hamiltonian will require an understanding of how a "good" system behaves. For this the transmission coefficient was introduced to compare the signal after the qubit to the signal before the qubit. While this method is applicable as it stems from both microwave engineering [11] and quantum mechanics [25] and it is agnostic of the drive pulse shape, the first results of using this metric are not physical. This indicates further consideration is required regarding exactly what part of the response signal is used, as well as the need to introduce normalization into this calculation.

An immediate next step for a cleaner signal for the transmission coefficient would be to extract the backpropagation from the signal on the top-edge such that all that is left is the top edge's response to the original drive pulse. Ideally, this should be done in such a way that is agnostic to the shape of the drive pulse. That way any possible changes to the pulse (e.g. using a non-Gaussian shape) would not impact such a method for quantifying the quality of the system. This would likely solve the problem of choosing a proper time-window over

which to integrate. The transmission coefficient, however, is meant to compare amplitudes, not necessarily the integral of amplitudes. Further consideration of how to properly normalize these integrals will enable better implementation of the transmission coefficient.

In parallel to determining a reliable metric, the architecture of the system itself should be considered in greater detail. While qubit interaction proves that it is possible to reverse the edge mode chirality (as is evident in the back-propagation), it is not studied how this manipulation would work. The qubits presented in this thesis were of a fixed frequency, meaning that their ground and excited states were hard-coded. It is possible and actually common place in quantum integrate circuits to use tunable qubits [14]. Such qubits' frequencies may be shifted when excited by an extra external flux that is only applied to the individual qubit in question. By modulating a qubit's frequency, it could be tuned in such a way that it does and does not interact with the lattice at the exact opportune moments. Whether both qubits or only one would need to be tunable for the final working system is to be determined. It is possible that both options are functional. The simplest, case, though, is to have one tunable qubit and one fixed.

To understand how changing the frequency of a single tunable qubit could flip the edge mode, it is suggested that only a single sublattice with one attached qubit is studied in detail. This means that a single Hofstadter model as presented in equation 54 may be modeled directly (with a qubit attached, of course) and the unitary transformation presented in equations 55 to 57 can be ignored. In this pre-transformation subsystem, it would be easier to study how the qubit could "absorb" and "inject" the signal in such a way, for example, as to flip the edge mode's chirality, which is equivalent to changing the Gaussian's phase θ by π , or enact an X gate. Such a sublattice simulation has been set up, but not yet studied in detail.

Another architecture choice is whether or not the lattice and control-qubit system should be able to easily manage several different single qubit gates. The alternative would be to design several of these systems each specifically optimized to enact one type of single qubit gate. The eventual processor made up of many of these lattices would then hard-code a particular set of gates. This decision between universal or application-specific focus is not of immediate concern, though.

With a proper metric and set architecture, the full Hamiltonian may then be optimized. Numerical methods such as automatic differentiation may prove incredibly efficient in determining optimal parameters for the pulse, qubit frequencies, and various couplings. Such a method, though, requires a reliable loss function (e.g. transmission coefficient). For example, in the case of an X gate, the transmission coefficient would ideally be 0 as the signal should be completely reflected into the opposite circular edge mode. In this case, automatic differentiation would run the simulation of the full system to find optimal Hamiltonian parameters that would bring the transmission coefficient as close to zero as possible. If an H gate was being studied instead, then the same method could be used, but with an ideal transmission coefficient of 0.5. While

the unoptimized one-qubit system demonstrated a backpropagation that one may consider a preliminary indication of a Hadamard gate, the qubit's oscillations did not land in the ground state. The parameters for this system could, however, serve as a good starting point for a Hamiltonian optimization process focused on tuning up a Hadamard gate. In the case of the H gate, GRAPE may also be implemented on the unoptimized system (for which the reflection was observed) such that it adjusts the pulse until the qubit ends in the ground state, i.e. $C = |0\rangle\langle 0|$ and the $\mathcal{H}_k$ is the terms of the full Hamiltonian related to the qubit. GRAPE, however, will only optimize the pulse for the sake of the qubit. Automatic differentiation could optimize many parameters for the sake of the entire system's response. The idea to use automatic differentiation for optimizing quantum systems has already been introduced for steady state systems [31] so could be extended to this use case.

Additionally, detailed characterization of the topological edge modes of such systems should be conducted. For example, one may study the energy as a function of the wave number, similar to as is present in figure 9. It should eventually be confirmed that there is indeed an energy band gap between the bulk and edge. The introductory study of the signal decay from the edge into the bulk presented in this thesis may be convincing enough to support continuing onwards with further detailed analysis of a lattice connected to qubit(s); however, as these systems become more defined, it will be helpful to have more thorough quantification of the topological features central to this proposed hardware's success. Another topological feature that may be considered is that expressed in equation 59. The assumption is that there is no coupling between the two edge modes. It could be possible, however, that there is some degree of interaction between the two system while still maintaining the overall time reversal symmetry of the system as a whole.

While the results of this thesis are preliminary with respect to the larger development of this new quantum hardware, studying the resonator responses on the edges and in the bulk, as well as the qubit responses in both one and two qubit cases, lays down the groundwork towards enacting single qubit gates on these systems. The results have not identified any major road blocks preventing the hypothesis that an LC resonator array with control qubit(s) could serve as a more robust qubit. Once better metrics have been established and the proposed Hamiltonian optimization techniques implemented, it is possible that single qubit gates may be realized on such a system.

Acknowledgements

I would like to thank my committee for their guidance. A particular thank you to Professor Eliška Greplova for reigniting my passion for research while encouraging my interests beyond the traditional hard sciences. Guliuxin Jin was incredible in answering my many questions regarding topological insulators, Overleaf, and plotting, as well as supporting me through derivations of which I embarked out of curiosity rather than necessity. Radoica Draškić introduced me to the scaffolding of this project. Anna (Ania) Dawid and Rodrigo Vargas Hernandez provided stimulating conversations regarding numerical methods for Hamiltonian optimization. André Melo and Anthony Polloreno graciously provided many helpful sanity checks. Dylan Everingham and Mirko Kemna, as well as the whole QMAI group created enlightening, enjoyable, and humane academic communities of which I am grateful to have been a part. My Berlin-based community lovingly kept me afloat the past two years. Ian Andrew Barber, Esq. has served as an excellent daily motivator and personal project manager for many years and only took the role more seriously as this project's deadlines approached. Thank you for always keeping me on track with a head held high and a spring in my step. Lastly, my family, who ingrained curiosity into my DNA. Thank you for all the love and support, even when I don't choose the easy way and especially when you don't understand what is it I work on, but wake up in the middle of the night to listen to me talk about it anyways.

This work is part of the project Engineered Topological Quantum Networks of the NWO Talent Programme Veni Science domain 2021 research program, which is financed by the Dutch Research Council (NWO) [5].

References

- [1] O. Dial. Eagle’s quantum performance progress, Mar 2022.
- [2] F. Arute, K. Arya, R. Babbush, et al. Quantum supremacy using a programmable superconducting processor. *Nature*, 574:505–510, 2019.
- [3] C. Nayak. Microsoft has demonstrated the underlying physics required to create a new kind of qubit, Mar 2022.
- [4] Microsoft Quantum Team. Developing a topological qubit, Sep 2018.
- [5] E. Greplová. Engineered topological quantum networks, March 2021.
- [6] R. Süssstrunk and S. D. Huber. Observation of phononic helical edge states in a mechanical topological insulator. *Science*, 349(6243):47–50, 2015.
- [7] N. Khaneja, T. Reiss, C. Kehlet, et al. Optimal control of coupled spin dynamics: design of NMR pulse sequences by gradient ascent algorithms. *Journal of Magnetic Resonance*, 172(2):296–305, 2005.
- [8] F. Motzoi, J. M. Gambetta, P. Rebentrost, and F. K. Wilhelm. Simple pulses for elimination of leakage in weakly nonlinear qubits. *Phys. Rev. Lett.*, 103:110501, Sep 2009.
- [9] M. A. Nielsen and I. L. Chuang. *Quantum Computation and Quantum Information*. Cambridge University Press, 10th anniversary edition, 2010.
- [10] S. M. Girvin. Circuit QED: Superconducting Qubits Coupled to Microwave Photons, 2014.
- [11] D. M. Pozar. *Microwave Engineering*. John Wiley & Sons, Inc., 4th edition, 2012.
- [12] D. Morin. *Introduction to Classical Mechanics, With Problems and Solutions*, chapter 15. Cambridge University Press, 2008.
- [13] C. R Hoffer. Superconducting qubit readout pulse optimization using deep reinforcement learning. Master’s thesis, Massachusetts Institute of Technology, Feb 2021.
- [14] D. I. Schuster. *Circuit Quantum Electrodynamics*. PhD thesis, Yale University, May 2007.
- [15] G. Jin and E. Grepova. Topological entanglement stabilization in superconducting quantum circuits, 2022.
- [16] C. L. Kane and E. J. Mele. Quantum spin hall effect in graphene. *Phys. Rev. Lett.*, 95:226801, Nov 2005.
- [17] A. Akhmerov, J. Sau, B. van Heck, et al. Topology in condensed matter: Tying quantum knots, 2021.

- [18] W. P. Su, J. R. Schrieffer, and A. J. Heeger. Solitons in polyacetylene. *Phys. Rev. Lett.*, 42:1698–1701, Jun 1979.
- [19] A. Pályi J. K. Asbóth, L. Oroszlány. *A Short Course on Topological Insulators*, volume 919 of *Lecture Notes in Physics*. Springer International Publishing, 2016.
- [20] J. Larson and T. Mavrogordatos. *The Jaynes–Cummings Model and Its Descendants*. IOP Publishing, Dec 2021.
- [21] R. Shankar. Topological insulators – a review, 2018.
- [22] S. Huber. Mechanical metamaterials, 2018.
- [23] B. Lian. Hofstadter butterfly: a momentum space method, and topological effects, Dec 2020.
- [24] P. de Fouquieres, S.G. Schirmer, S.J. Glaser, et al. Second order gradient ascent pulse engineering. *Journal of Magnetic Resonance*, 212(2):412–417, 2011.
- [25] D. J. Griffiths. *Introduction to quantum mechanics*. Prentice Hall, 1995.
- [26] M. Kjaergaard, M. E. Schwartz, J. Braumüller, et al. Superconducting qubits: Current state of play. *Annual Review of Condensed Matter Physics*, 11(1):369–395, Mar 2020.
- [27] C. C. Gerry and P. L. Knight. *Introductory Quantum Optics*. Cambridge University Press, 2005.
- [28] A. Wipf. *Statistical Approach to quantum field theory: An introduction*, volume 864 of *Lecture Notes in Physics*. Springer Verlag, 2013.
- [29] W. Weckesser. odeintw, Jun 2019.
- [30] D. M. Pozar. *Microwave Engineering*. John Wiley & Sons, Inc., 4th edition, 2012.
- [31] R. A. Vargas-Hernández, R. T. Chen, K. A. Jung, et al. Fully differentiable optimization protocols for non-equilibrium steady states. *New Journal of Physics*, 23(12):123006, Dec 2021.