



UNIVERSITÀ
DI PAVIA

UNIVERSITÀ DEGLI STUDI DI PAVIA

Dipartimento di Fisica “Alessandro Volta”

Corso di Laurea Magistrale in Scienze Fisiche

High-probability two-qubit gate in nonlinear quantum walks

Tesi per la Laurea Magistrale di:
Matilde MICCOLI

Relatore:
Prof. Dario GERACE

Correlatore:
Dr. Francesco SCALA

Anno accademico 2025 – 2026

Contents

Introduction	2
1 Theoretical background	5
1.1 Introduction to Quantum Computing Theory	5
1.1.1 Qubits encoding	5
1.1.2 Elementary quantum gates	7
1.2 Quantum computing platforms	10
1.2.1 DiVincenzo Criteria	11
1.3 Photonic Quantum Computing	13
1.3.1 Single-photon qubit encoding	14
1.3.2 The Kerr non-linearity	15
1.3.3 Linear optical quantum computing: the KLM paradigm	18
1.3.4 State of the art in Photonic Quantum Computing	24
1.4 Alternative quantum computing platforms	30
1.5 Measurement-based quantum computation	32
Summary	33
2 Two-qubit gates in linear optical quantum computing	34
2.1 The KLM protocol	35
2.2 Theoretical bounds for post-selected gates	38
2.2.1 Improving the CZ gate	38
2.2.2 Quantitative bounds	39
2.3 Implementations of linear-optics CNOT	41
2.3.1 CNOT via quantum erasure	41
2.3.2 CNOT with a single ancilla photon	43
2.4 Probabilistic CNOT without ancilla photons	45
2.4.1 Experimental realizations: CNOT without ancilla photons	50
Summary	54
3 Nonlinear quantum walks and interferometers	55
3.1 Quantum walks and probabilistic two-qubit gate	55
3.1.1 Probabilistic CNOT gate in linear quantum walks	56
3.1.2 Experimental realization with non-interacting QWs.	62
3.2 Quantum nonlinear interferometers and...	64
3.2.1 Deterministic entangling gates in Kerr-nonlinear interferometers	66
Summary	71

4	Algorithmic optimization of photonic two-qubit gates	73
4.1	The nonlinear quantum walk	74
4.1.1	Hamiltonian and waveguide lattice	74
4.1.2	Qubit encoding	76
4.2	Definitions of Success probability and Fidelity	77
4.2.1	Success probability	77
4.2.2	Two-qubit gate fidelity	78
4.3	Optimization strategy	80
4.3.1	Parameters, search space and cost function	80
4.3.2	Optimization algorithm and strategy	81
4.3.3	Multi-run aggregation	82
4.4	Numerical results: CNOT in linear regime	83
4.5	CNOT gate in nonlinear quantum walks	90
4.5.1	Weak non-linearities	91
4.5.2	Strong non-linearities	92
4.5.3	Focus on nonlinear results for $\Gamma = 1$ ($\Gamma/J_{\max} = 0.125$)	93
4.6	Physical plausibility of optimized parameters	97
	Summary	98
	Conclusions	99
A	Simulation code	102
A.1	Single-run optimization script	102
A.2	Multi-run aggregation script	107
B	Numerical results for small non-linearity values	112
B.1	Non-linear regime of $\Gamma = 0.1$ ($\Gamma/J_{\max} = 0.020$)	112
B.2	Non-linear regime of $\Gamma = 0.25$ ($\Gamma/J_{\max} = 0.045$)	114
B.3	Non-linear regime of $\Gamma = 0.5$ ($\Gamma/J_{\max} = 0.083$)	115
B.4	Non-linear regime of $\Gamma = 0.75$ ($\Gamma/J_{\max} = 0.107$)	119
	Bibliography	125

Introduction

In the long term, quantum computers promise a real advantage over classical computing machines for problems such as cryptography, drug design, and machine learning. This advantage comes from genuinely quantum resources, such as superposition, interference, and entanglement, which have no classical counterpart. The basic unit of quantum information is the qubit (quantum bit of information), i.e., essentially a two-level quantum system whose state can be visualized as a well defined point on the Bloch sphere (see, e.g., Fig. 1.1 in the following Chapter) [1]. Any quantum algorithm can be built from a small, universal set of elementary gates acting on one or two qubits at a time; among these, entangling two-qubit gates such as the controlled-NOT (CNOT) are essential, since no sequence of single-qubit operations alone can generate entanglement between separate qubits.

Turning this abstract circuit model into a physical machine, however, is a hard task. A good qubit platform must satisfy several conditions at once, informally summarized by the so called “DiVincenzo criteria” [2]: 1) qubits must be well characterized and 2) arbitrarily initialized; 3) they must keep their quantum coherence for much longer than the time needed to perform a single quantum gate; 4) a universal gate set must be implementable on the qubits register; 5) the result of a computation must be measurable on each qubit individually.

These requirements pull in opposite directions: a qubit must be well isolated from its environment to preserve coherence, yet accessible enough to be controlled and read out. Different physical platforms strike this balance very differently, from nuclear and electron spins to trapped ions and superconducting circuits, each with its own characteristic coherence time and gate speed.

Among these platforms, photonic qubits occupy a distinctive position. Photons interact only weakly with their environment, so, unlike matter-based qubits, they do not really decohere during propagation: information is instead lost mainly through photon loss in the optical components themselves. Photons (in the near-infrared band) are also directly compatible with the optical-fibre infrastructure already used for classical communication, making them a natural carrier for both quantum computation and quantum communication on the same physical platform.

This same weak interaction with the environment is also the main problem, preventing photonic qubits to be widespread among the possible candidates to practical quantum information processing, yet. In fact, photons barely interact with each other, too. Building a two-qubit gate, i.e., the basic tool needed for universal quantum computation, requires some kind of photon-photon interaction. No natural non-linearity is strong enough to do this at the level of single photons [3, 4]. Linear optical quantum computing (LOQC), proposed by Knill, Laflamme and Milburn [5], gets around this

problem instead of solving it directly: two-qubit gates are built using only linear optical elements, extra ancilla photons, and measurements that act as an effective, “hidden” non-linearity. This paradigm works, but only in a probabilistic way. For a CNOT gate, the most common two-qubit gate to be implemented, the maximum success probability is fixed at $1/9$ when no ancilla photons and no feed-forward are used [6, 7, 8]. This bound was reached by beam-splitter schemes [9, 10] and confirmed experimentally, both with bulk optics [11, 12, 13] and with integrated silica-on-silicon chips [14]. Alternative, measurement-based paradigms such as cluster-state computation offer a different route to the same probabilistic gates, shifting the difficulty from building deterministic entangling gates to preparing a large entangled resource state beforehand, at the cost of a substantial photon-number overhead.

More recently, a simpler way to reach the same $1/9$ bound was theoretically proposed (and then experimentally realized). Instead of engineering many beam splitters, the two logical qubits are encoded in the position of two identical photons moving through a line of coupled waveguides: this is what is defined as a one-dimensional *quantum walk*. Y. Lahini et al. [15] have theoretically shown a few years ago that the natural two-photon interference in such a one-dimensional waveguide lattice reproduces the optimal linear-optical CNOT gate, using only six waveguides and no active control at all. This result was later confirmed by experiments, after realization of the very same quantum walk design in an integrated photonic platform [16], showing that the idea is robust and not tied to one specific technology.

At the same time, a second and independent line of research asked a different question: what happens if we add some photon-photon interaction? Within this context, a work from the Department of Physics at the University of Pavia from F. Scala et al. [17] has recently put forward the idea that a distributed self-Kerr non-linearity, spread across a complex photonic interferometer, theoretically removes the $1/9$ bound completely. In this proposed scheme, the CNOT gate becomes fully deterministic, with fidelity as close to 100% as desired, and no leakage of excitations out of the computational basis states, whenever the system is initialized into one of them (or any arbitrary superposition). The price to pay is a more complex interferometer design, built from several concatenated blocks, which seems currently out of reach from realization in a practical device. In parallel, a quite similar picture has emerged in a related context: quantum optical neural networks built from trainable linear nets and single-site Kerr non-linearities have been shown to vastly outperform linear-optical bounds even when the non-linearity is weak or imperfect, well below the ideal value [18, 19].

Now, the contribution of the present thesis sits between these two ideas. We ask a simple question: can we keep the architectural simplicity of the quantum-walk CNOT gate introduced by Y. Lahini et al., while still using a distributed self-Kerr non-linearity to beat the theoretical $1/9$ bound imposed by LOQC principles? And if so, how does the gate behave as we move continuously from the purely linear case to a strongly non-linear one? To answer these meaningful and still open questions, we study a system of six straight, evanescently coupled waveguides, described by an Hamiltonian with a uniform self-Kerr non-linearity added at each site. We numerically optimize the couplings between waveguides for several fixed values of the non-linearity strength, and we track how the trade-off between fidelity and success probability changes as the non-linearity is switched on and increased. As we will show, this simple picture already captures a smooth crossover between the two regimes described above: a purely probabilistic

gate at zero non-linearity, and a regime where fidelity and success probability can both be made simultaneously large as the non-linearity is increased. This results, which is the main outcome of the work, opens up new avenues for the optimization of two-qubit gates in one-dimensional quantum walks, largely overcoming the difficulties of the KLM paradigm, and thus constitutes a major theoretical result per se.

The presentation of this thesis is organized as follows.

First, Chapter 1 introduces the basic theory of quantum computation, from the circuit model and elementary gates to the physical requirements a good platform must satisfy, highlighting the DiVincenzo criteria, and then focuses on the photonic approach: single-photon qubit encoding, the role of Kerr-type non-linearities, the KLM protocol, and a brief overview of the current state of the art in photonic hardware and of alternative platforms and measurement-based paradigms.

Then, in Chapter 2 the theoretical bounds on postselected linear-optical gates are worked out in detail, by explicitly presenting the specific six-waveguide CNOT construction of Ralph et al. [10], which is the linear-optical benchmark used throughout this work.

Chapter 3 introduces the already mentioned quantum-walk picture of Y. Lahini et al., as well as the deterministic non-linear interferometer of Scala et al., the two results this thesis tries to connect.

Chapter 4 extensively presents the original contribution of this work: the physical model, the optimization strategy, and the numerical results obtained for increasing values of the non-linearity, compared against both the linear-optical bound and the fully deterministic non-linear gate. The main (and original) Python codes developed for the extensive numerical work performed during the thesis are reported in the Appendices.

Chapter 1

Theoretical background

1.1 Introduction to Quantum Computing Theory

Quantum computation and quantum information deal with the study of the processing tasks that can be accomplished using quantum mechanical systems. The emergence of quantum information processing, as a theoretical framework, has sparked a wealth of novel ideas, in a very practical sense, regarding the design of physical computing systems that work by exploiting the use of quantum mechanical principles. The ongoing aim is the development of operational laboratory systems capable of carrying out this computational model and a new paradigm of information processing.

In this first chapter, we aim to introduce the reader to the fundamentals of quantum computing theory, and the basic quantum mechanical tools that will be used in the following. To do so, we will follow the structure of the standard textbook reference by Nielsen and Chuang, *Quantum Computation and Quantum Information*([1]).

1.1.1 Qubits encoding

The bit is the fundamental concept of classical computation and classical information. Quantum computation and quantum information are built upon an analogous concept, the quantum bit, also called the *qubit*. Just as a classical bit can be in one of two possible states (i.e., either 0 or 1), a qubit is also characterized by a pair of quantum states, which can be $|0\rangle$, $|1\rangle$ or a linear combinations of these two, i.e., the generic superposition:

$$|\Psi\rangle = \alpha |0\rangle + \beta |1\rangle . \quad (1.1)$$

The coefficients α and β are complex numbers. Qubits can be understood as unitary vectors in a two-dimensional complex vector space. The special states $|0\rangle$ and $|1\rangle$ are defined as the computational basis states, and form an orthonormal basis for this vector space.

When we measure a qubit we get either the result 0, with probability $|\alpha|^2$, or the result 1, with probability $|\beta|^2$. It is important to remind, as a fundamental quantum mechanical principle, that the measurement process changes the state of the qubit, by collapsing it from its superposition of $|0\rangle$ and $|1\rangle$ to the specific state consistent with the measurement result.

Naturally, $|\alpha|^2 + |\beta|^2 = 1$, since the probabilities must sum to one. Geometrically, we

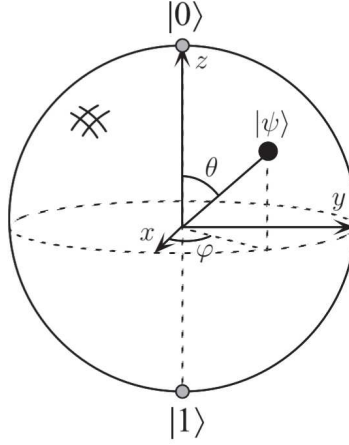


Figure 1.1: Bloch sphere representation of a single qubit quantum state. Picture taken from Ref. [1].

can interpret this as the condition that the qubit state be normalized to length 1. Hence, we can rewrite 1.1, up to a phase factor, as:

$$|\Psi\rangle = \cos \frac{\theta}{2} |0\rangle + e^{i\phi} \sin \frac{\theta}{2} |1\rangle, \quad (1.2)$$

which is now parametrized by two real numbers, θ and ϕ , uniquely defining a single point on the three-dimensional sphere with unit radius, called the *Bloch sphere*, as schematically represented in Fig. 1.1.

We now suppose to have two qubits, whose Hilbert space is then spanned by four computational basis states, typically denoted $\{|00\rangle, |01\rangle, |10\rangle, |11\rangle\}$. The generic quantum state of a pair of qubits can then exist in a given superposition of these four states, i.e.:

$$|\Psi\rangle = \alpha_{00} |00\rangle + \alpha_{01} |01\rangle + \alpha_{10} |10\rangle + \alpha_{11} |11\rangle \quad (1.3)$$

The measurement result, i.e., $x \in \{00, 01, 10, 11\}$, occurs with probability $|\alpha_x|^2$, with the state of the qubits after the measurement being defined with certainty as $|x\rangle$. Furthermore, the condition that probabilities sum to one is now expressed by the *normalization condition*:

$$\sum_{x \in \{0,1\}} |\alpha_x|^2 = 1$$

The states $|\Psi\rangle$ can be *entangled*, meaning that they cannot be written as a tensor product of the states of the two individual qubits. We remind that the tensor product $V \otimes W$ of two vector spaces V and W is the vector space spanned by elements $v \otimes w$ with $\dim(V \otimes W) = \dim(V) \dim(W)$, which satisfies the properties $(v_1 + v_2) \otimes w = v_1 \otimes w + v_2 \otimes w$ and $\alpha(v \otimes w) = (\alpha v) \otimes w = v \otimes (\alpha w)$.

Generalizing the concepts above, we notice that the general state of n qubits is specified by a 2^n -dimensional complex vector. Quantum mechanics can be expressed in two mathematically equivalent ways. The first relies on state vectors, while the second makes use of a tool called the density operator (or density matrix). Although both approaches lead to the same results, the density operator formulation is often more convenient and elegant to describe certain situations that are commonly encountered in

quantum mechanics. More precisely, suppose a quantum system is in one of a number of states $|\psi_i\rangle$, where i is an index, with respective probabilities p_i . We identify $\{p_i, |\psi_i\rangle\}$ as an *ensemble of pure states*. The density operator (or density matrix) for the system is defined by the equation:

$$\rho \equiv \sum_i p_i |\psi_i\rangle \langle \psi_i|$$

A density operator, defined like that, is characterized by the following conditions:

- **Trace condition:** ρ has trace equal to 1.
- **Positivity condition:** ρ is a positive operator.

From this view of quantum mechanics we can write a formulation of the probability of a specific result. Quantum measurements are described by a collection $\{M_m\}$ of measurement operators, also known as Kraus operators¹. These are operators acting on the state space of the system being measured. The index m refers to the measurement outcomes that may occur in the experiment. If the state of the quantum system is ρ immediately before the measurement then the probability that result m occurs is given by:

$$p(m) = \text{Tr}(M_m^\dagger M_m \rho) \quad (1.4)$$

The measurement operators satisfy the completeness equation:

$$\sum_m M_m^\dagger M_m = I \quad (1.5)$$

The density operator approach to quantum mechanics is really useful for two applications: the description of quantum systems whose state is not known, and the description of subsystems of a composite quantum system.

1.1.2 Elementary quantum gates

In Quantum Computing theory, a quantum logic gate is a basic quantum circuit operating on a small number of qubits. Quantum logic gates are the building blocks of quantum circuits. Let us start from single qubit gates. It is possible to represent these gates by a matrix, but only if the matrix U is unitary, i.e., satisfying $U^\dagger U = I$. The latter is the only constraint on quantum gates. In fact, any unitary matrix acting on a 2^n dimensional Hilbert space defines a valid quantum gate acting on a system of n qubits.

Here we list the most commonly employed single qubit gates:

- *Identity gate*(I): is basis independent and does not modify the quantum state. The identity matrix is defined for a single qubit as

$$I = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$

¹Another frequently used notation identifies $O_m = M_m^\dagger M_m$ as a measurement operator, but we will not refer to this.

- *Pauli gates* (X, Y, Z): are discrete rotations around the x, y, z axes of the Bloch sphere by π radians. The Pauli matrix ($\sigma_x, \sigma_y, \sigma_z$) are defined for a single qubit as

$$\sigma_x = X = NOT = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \quad \sigma_y = Y = \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} \quad \sigma_z = Z = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix}$$

- *The Hadamard gate* (H): is just a rotation of the sphere about the $\hat{y}$ axis by 90° , followed by a rotation about the $\hat{x}$ axis by 180° . The Hadamard matrix is defined for a single qubit as

$$H = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}$$

- *The phase shift*: is a family of single-qubit gates that map the basis states

$$|0\rangle \rightarrow |0\rangle \quad |1\rangle \rightarrow e^{i\phi} |1\rangle$$

The probability of measuring $|0\rangle$ or $|1\rangle$ is unchanged after applying this gate, however it modifies the phase of the quantum state. The most important gates from this family are:

- $\pi/8$ gate (T):

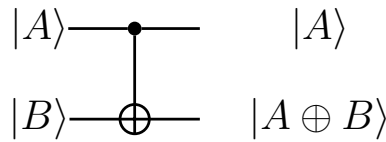
$$T = \begin{bmatrix} 1 & 0 \\ 0 & e^{i\pi/4} \end{bmatrix} = e^{i\pi/8} \begin{bmatrix} e^{-i\pi/8} & 0 \\ 0 & e^{i\pi/8} \end{bmatrix}$$

- *phase gate* (S):

$$S = \begin{bmatrix} 1 & 0 \\ 0 & i \end{bmatrix}$$

In addition to single-qubit operations, there are a few elementary multi-qubit operations that are essential. In fact, it is possible to show that any quantum computation acting on an arbitrary number of qubits can be implemented through a finite set of gates, known as a *universal set for quantum computation* [1]. In particular, a universal set is given by arbitrary two qubits rotations and at least one of the following entangling multi-qubit gates.

- *The CNOT gate*: The controlled-NOT or CNOT gate acts between two qubits, known as the *control* and the *target* qubit, respectively. The circuit representation is shown in the following figure:



The top line represents the control qubit, while the bottom line represents the target qubit. Where $\oplus$ represent a bit-wise sum.

We can describe the action of the gate as follows: if the control qubit is set to 0, then the target qubit is left unchanged; if the control qubit is set to 1, then the target qubit is flipped. In equations:

$$|00\rangle \rightarrow |00\rangle \quad |01\rangle \rightarrow |01\rangle \quad |10\rangle \rightarrow |11\rangle \quad |11\rangle \rightarrow |10\rangle$$

The CNOT matrix is

$$U_{CN} = CNOT = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{bmatrix} \quad (1.6)$$

The requirement that probability be conserved is expressed in the fact that U_{CN} is a unitary matrix, that is, $U_{CN}^\dagger U_{CN} = I$.

The CNOT gate is capable of generating entanglement between qubits, but this does not happen unconditionally: it depends on the input state. The key is to have the control qubit in a superposition. For instance, applying a Hadamard gate to the control qubit and then a CNOT with the state initialised in $|00\rangle$ yields one of the Bell states:

$$CNOT(H \otimes I) |00\rangle = \frac{1}{\sqrt{2}}(|00\rangle + |11\rangle)$$

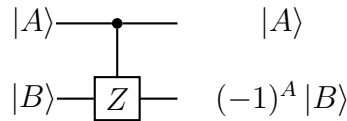
which cannot be written as a tensor product of two individual qubit states, which is the definition of entangled.

This works because the CNOT introduces a conditional correlation and when the control is in superposition, this conditional logic extends across the entire superposition, binding the two qubits in an inseparable way. No sequence of single-qubit gates can achieve this, since they act on each qubit independently and are therefore unable to create correlations between distinct qubits.

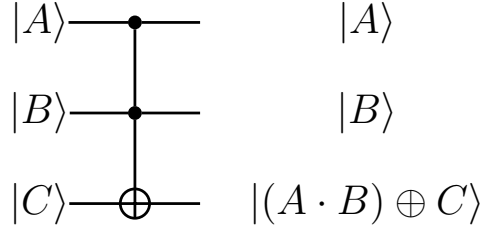
- *The CZ gate:* This gate flips the phase of the target qubit if the control qubit is in the $|1\rangle$ states. The CZ matrix is

$$U_{CZ} = CZ = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & -1 \end{bmatrix} = (I \otimes H)CNOT(I \otimes H) \quad (1.7)$$

The circuital representation is:



- *The Toffoli gate:* The Toffoli gate, also known as the controlled-controlled-NOT (CCNOT) gate, is a three-qubit gate with two control qubits and one target qubit. The circuit representation is shown below:



The gate flips the target qubit if and only if both control qubits are set to $|1\rangle$; otherwise, all qubits are left unchanged. In equations:

$$\begin{array}{llll} |000\rangle \rightarrow |000\rangle & |001\rangle \rightarrow |001\rangle & |010\rangle \rightarrow |010\rangle & |011\rangle \rightarrow |011\rangle \\ |100\rangle \rightarrow |100\rangle & |101\rangle \rightarrow |101\rangle & |110\rangle \rightarrow |111\rangle & |111\rangle \rightarrow |110\rangle \end{array}$$

The Toffoli matrix, acting on the computational basis ordered as

$$\{|000\rangle, |001\rangle, |010\rangle, |011\rangle, |100\rangle, |101\rangle, |110\rangle, |111\rangle\}$$

is:

$$U_{CCN} = \text{Toffoli} = \begin{bmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \end{bmatrix}$$

As with the CNOT and CZ gates, U_{CCN} is unitary, satisfying $U_{CCN}^\dagger U_{CCN} = I$.

The Toffoli gate is the quantum analog of the classical AND gate (with an ancilla qubit), making it a key primitive for implementing classical logic reversibly within a quantum circuit.

1.2 Quantum computing platforms

Building a quantum computer requires not only finding a stable physical substrate for qubits — i.e., one that preserves their quantum nature — but also identifying a system in which their evolution can be precisely controlled. In addition, one needs the ability to initialize qubits in well-defined starting configurations, and to extract the final state of the system through measurement.

Hence, a fundamental tension lies at the heart of quantum computing: on one side the system must be sufficiently shielded from its environment to maintain quantum coherence, on the other, though, its qubits must remain accessible enough to be manipulated during computation and final read out.

In this respect, a central figure of merit when evaluating any candidate implementation is quantum noise, also referred to as decoherence. The main reason is that the maximum feasible length of a quantum computation is approximately determined by the ratio between τ_Q , the timescale over which the system sustains quantum-mechanical

System	τ_Q	τ_{op}	$n_{op} = \tau_Q/\tau_{op}$
Nuclear spin (P in Si)	$10^0\text{--}10^2$ s	$10^{-6}\text{--}10^{-5}$ s	$10^7\text{--}10^9$
Electron spin (Si quantum dot)	$10^{-1}\text{--}10^1$ s	$10^{-7}\text{--}10^{-6}$ s	$10^6\text{--}10^8$
Ion trap	$10^1\text{--}10^3$ s	$10^{-6}\text{--}10^{-4}$ s	$10^5\text{--}10^9$
Superconducting (transmon)	$10^{-4}\text{--}10^{-3}$ s	$10^{-8}\text{--}10^{-7}$ s	$10^4\text{--}10^5$
Photonic (linear optics)	$\gg \tau_{op}$	$10^{-15}\text{--}10^{-12}$ s	$\gg 10^9$

Table 1.1: Decoherence times τ_Q , operation times τ_{op} , and estimated number of coherent operations n_{op} for the main quantum computing platforms as of 2025. Values for nuclear spins from [20]; electron spins from [21]; ions in trapped potentials from [22]; superconducting transmon qubits from [23]; photonic qubits from [24].

coherence, and τ_{op} , the time required to carry out elementary unitary operations, which must involve at least two qubits. In many physical systems these two timescales are not independent, as they are both governed by how strongly the system couples to its surroundings. Nonetheless, the ratio $n_{op} = \tau_Q/\tau_{op}$ turns out to span a remarkably broad range of values across different platforms. A brief summary is provided in Table 1.1, with values taken from the recent literature.

Among all the different platforms, we specifically notice that the photonic platform requires special consideration. In fact, single photons do not undergo decoherence in the traditional sense, because they do not interact with a thermal environment during propagation. This is a potentially key advantage of photonic qubits with respect to other degrees of freedom. In this case, τ_Q is not limited by coupling to external degrees of freedom, but rather by photon loss in the optical components (such as, e.g., waveguides, beam splitters, single-photon detectors), which sets an effective fidelity ceiling on the computation. The operation time, τ_{op} , is instead set by the speed of light through the circuit, making n_{op} formally very large. For this reason, the figure of merit that best characterizes photonic platforms is not the ratio τ_Q/τ_{op} but rather the total transmission efficiency of the optical circuit, sometimes defined η [25, 5].

1.2.1 DiVincenzo Criteria

To outline which types of physical platforms are promising for practical quantum information processing, it is necessary to briefly mention the famous set of *DiVincenzo Criteria* [2]. These are five fundamental requirements formulated by physicist David DiVincenzo in 2000, which define the necessary conditions for building a functional quantum computer. Moreover, there are two additional Criteria for Quantum Communication (i.e., meant to transfer quantum information between different computers). Here is the list of criteria as originally formulated:

1. Well-defined qubits: A scalable physical system containing a collection of qubits is needed, in particular, the quantum system must have well-characterized qubits. To say a qubit is 'well-characterized' means we have a complete and precise map of how it behaves. Specifically, we need to know its internal energy levels (which define the $|0\rangle$ and $|1\rangle$ states), whether it accidentally interacts with other nearby states or qubits and exactly how it responds to the external signals we use to control it.

2. Qubit initialization: The ability to initialize the state of the qubits to a simple reference state. This is necessary for two reasons. First, you need to set your system to a starting value (like $|0\rangle$) before you begin any calculation. Second, for quantum error correction to work, you need a steady stream of 'fresh' qubits in a clean state (like the $|0\rangle$ state) to keep the process running smoothly. There are two main ways to reset qubits to a starting state. You can either cool the system down until it naturally settles into its lowest energy state (the ground state), or you can measure the qubit forcing it into a specific state.

3. Long coherence times: Quantum coherence times need to be substantially longer than gate operation times. Qubits must maintain their quantum state long enough to complete calculations before decoherence occurs. Decoherence is very dangerous for quantum computing, since if it acts for very long, the capability of the quantum computer will not be so different from that of a classical machine.

4. Universal gates set: A “universal” set of quantum gates that can implement any quantum algorithm through their combinations. In simple terms, this means having the tools to manipulate qubits individually and make them interact in pairs (like the CNOT gate). If you have these basic building blocks, you can perform any complex quantum algorithm.

5. Qubit-specific measurement: The result of a computation must be read out, so it is necessary to have the ability to measure specific qubits without disturbing others.

6. Interconversion between stationary and flying qubits: The ability to convert between qubits used for computation and qubits used for transmitting quantum information.

7. Faithful transmission of flying qubits: The ability to transmit qubits between specific locations with high fidelity. This means ensuring that the quantum information stays intact as it travels.

Universal gates set for Quantum Computation

As already outlined in the 4th DiVincenzo criterion, the existence of a universal set of elementary quantum gates is a fundamental requirement for Quantum Computation (QC), as it ensures that the hardware can execute any possible quantum algorithm.

To understand the importance of this concept, and specifically the crucial role of the CNOT gate, we must start with a key decomposition theorem: *any unitary operator acting on a d -dimensional Hilbert space can be written as a product of two-level unitary matrices*. These are operations that act non-trivially only on a two-dimensional subspace spanned by two computational basis states, leaving all others unchanged.

Building on this, it can be proven that any such two-level unitary operation on the state space of n qubits can be implemented exactly using a combination of single-qubit gates and the CNOT gate. This highlights the indispensable role of the CNOT: it provides the entangling power needed to couple different qubits, a task that no single-qubit gate

can achieve on its own.

The final ingredient concerns the single-qubit gates themselves. A fundamental result, known as *the Z – Y decomposition*, guarantees that any unitary operation U on a single qubit can be written, up to a global phase, as a sequence of rotations around the axes of the Bloch sphere:

$$U(\theta, \phi, \lambda) = \begin{bmatrix} \cos(\theta/2) & -\sin(\theta/2)e^{i\lambda} \\ \sin(\theta/2)e^{i\phi} & \cos(\theta/2)e^{i(\lambda+\phi)} \end{bmatrix} = R_z(\phi)R_x(\theta)R_z(\lambda)$$

This means that single-qubit gates are entirely characterized by continuous rotation angles, defined by the equation $R_{\hat{n}}(\theta) = \exp\left(-i\frac{\theta}{2}\hat{n} \cdot \vec{\sigma}\right)$, where $\hat{n}$ is a real unit vector representing the rotation axis and $\vec{\sigma}$ is the vector of the Pauli's matrices, that is:

$$\hat{n} \cdot \vec{\sigma} = n_x\sigma_x + n_y\sigma_y + n_z\sigma_z$$

However, in practical physical realizations, implementing exact continuous angles is technically challenging and prone to noise. To address this, another theorem guarantees that any such single-qubit unitary can be approximated to arbitrary accuracy using only a finite and discrete set of gates, most commonly the Hadamard gate (H), the phase gate (S), and the $\pi/8$ gate (T).

By concatenating these results together, we can move from an abstract unitary operation to a sequence of discrete, implementable gates: $\{H, S, T, \text{CNOT}\}$. This transition to a discrete gate set has profound practical implications for quantum error correction. As DiVincenzo notes, working with a limited repertoire of operations allows errors to be discretized and corrected, which is essential for building a fault-tolerant quantum computer. In this framework, the CNOT gate remains the vital link; it is the only element capable of generating the entanglement required to make quantum computation truly universal and fundamentally more powerful than its classical counterpart.

1.3 Photonic Quantum Computing

It is well established that quantum computers can be built using a variety of platforms, each relying on different underlying physical systems. Among the most promising approaches, optical quantum systems are attractive ones for quantum computing, because of their inherent quantum properties and the minimal decoherence exhibited by photons, as already mentioned. Therefore, the aim of this section is to introduce the field of photonic quantum computing (PQC). We will rely on [1], as elsewhere in this Chapter, but we will also make use of material reported in Ref. [26].

Photons are chargeless particles, they can be guided through long propagation distances with low loss in optical fibers, efficiently delayed by using phase shifters, and easily combined by using directional couplers and beam splitters. Their resistance to decoherence tends to keep the stored quantum information stable.

However, photons do not naturally interact with each other, or even with most matter. This is the main drawback preventing photonic qubits to be immediately of interest to realize quantum computing platforms. In fact, sizeable photon-photon interactions would be essential for implementing two-qubit quantum gates. In principle, photons can be made to interact with each other by using non-linear optical media which mediate

interactions. Typically, though, such non-linearities are too weak to be of practical relevance to this purpose, for what we have been knowing so far. Here we will expose the basic principles of optical quantum computing, without being worried about the way sizeable photon-photon interactions are actually produced, which will be mentioned in the following.

1.3.1 Single-photon qubit encoding

In optical quantum computing, qubits can be encoded by exploiting either polarization or spatial degrees of freedom [27]:

- In the case of a *polarization qubit*, the two modes correspond to the internal polarization state of the photon, where

$$|0\rangle_P = |H\rangle \quad |1\rangle_P = |V\rangle$$

with only a single spatial mode being involved.

- For a spatial qubit, the *single-rail* encoding represents the qubit as

$$|1\rangle_{sr} = a^\dagger |0\rangle \quad |0\rangle_{sr} = a |1\rangle$$

respectively the first one represents a photon that occupies the waveguide, the second one represents the waveguide contains no photon.

- Spatial modes can further be employed to define the so-called *dual-rail qubit*, where the logical states are encoded in a single photon distributed across two distinct waveguides:

$$|0\rangle_{dr} = a_1^\dagger |\Omega\rangle = |1\rangle \otimes |0\rangle = |1, 0\rangle \quad |1\rangle_{dr} = a_2^\dagger |\Omega\rangle = |0\rangle \otimes |1\rangle = |0, 1\rangle$$

where $|\Omega\rangle = |0, 0\rangle$ is the vacuum state.

- There is another way to encode a qubit, which is putting a single photon in a superposition state of two time slots, called time bins. This is the so called *time-bin encoding*, which is mostly used for quantum communications applications, in particular for quantum cryptography protocols. To encode such type of qubit, you have to use a very short pulsed excitation laser, pulse a single photon into a standard Mach-Zehnder interferometer or any optical element that will increase the path length of the light, and then you have two possible situations: the optical pulse will be ahead (the so-called late L state) or behind time (the so-called early E state). The time-bin superposition state can thus be written:

$$|\psi\rangle_{time-bin} = \frac{1}{\sqrt{2}}(|L\rangle + |E\rangle)$$

The dual-rail and polarization encodings are mathematically equivalent representations, and it is possible to convert between them experimentally by means of polarizing beam splitters.

In the single-rail encoding, implementing single-qubit gates is not feasible, as it is impossible to act on the “no photon” state, namely the vacuum of the electromagnetic field. However, two-qubit entangling gates can be realized in a straightforward manner using a simple beam splitter. In contrast, the dual-rail encoding naturally allows for the implementation of single-qubit operations, while two-qubit gates become significantly more challenging to realize.

Before discussing the optical elements used in practice, it is useful to introduce the quantum states most commonly encountered in photonic implementations.

The *Fock states*, also known as number states, are quantum states that describe a definite number of photons in a given system. They are eigenstates of the number operator, $n = a^\dagger a$ (with $a^\dagger$ and a creation and annihilation operators of single mode quanta, respectively). Mathematically, the Fock state $|n\rangle$, satisfies:

$$a^\dagger a |n\rangle = n |n\rangle$$

where n represents the number of photons.

In order to encode qubits by using single-photon Fock states, it is necessary to efficiently create single photons. A possible way is by attenuating the output of a laser. In fact, a laser emits *coherent states* with a very good approximation. The latter are eigenstates of the destruction operator, which are mathematically defined as:

$$|\alpha\rangle = e^{-\frac{|\alpha|^2}{2}} \sum_{n=0}^{\infty} \frac{\alpha^n}{\sqrt{n!}} |n\rangle$$

where $|n\rangle$ is a n -photon energy eigenstate. The interesting feature that makes this state useful for our purpose is that an attenuated coherent state is still a coherent state, but weaker, and a weak coherent state can be made to have just one photon, with high probability.

On the other hand, single quantum emitters, such as semiconductor quantum dots, are able to emit highly pure and indistinguishable single-photons [27]. For the rest of this thesis, we will assume that single-photon Fock states can be produced and be used as the source of photonic qubits encoding.

1.3.2 The Kerr non-linearity

As discussed above, two-qubit gates are the central challenge in photonic quantum computing: unlike single-qubit operations, which are readily implemented with linear optical elements such as beam splitters and phase shifters, as shown below, entangling gates require some form of photon–photon interaction. A natural candidate for mediating such interactions is provided by non-linear optical media, and in particular by the *optical Kerr effect*.

This is a third-order non-linear optical phenomenon (i.e., originating from some $\chi^{(3)}$ susceptibility of an optically dense medium) in which the refractive index n of the medium in which photons propagate acquires an intensity-dependent contribution, such as

$$n = n_0 + n_2 I, \tag{1.8}$$

in which n_0 is the linear refractive index, n_2 is the non-linear Kerr coefficient, and I is the local optical intensity, where optical intensity is defined as the power per unit

area of a light wave. A detailed discussion on the origin of this non-linearity, together with its full quantum-optical treatment, is deferred to Chapter 3 (in particular, see Sec. 3.2); here we restrict ourselves to the minimal formalism needed to introduce the basic principles of optical quantum computing.

In fact, the Kerr effect manifests itself into two distinct regimes: cross-phase modulation (XPM) and self-phase modulation (SPM). In the *cross-phase modulation* regime, the presence of one photon in one mode induces a phase shift on a second photon propagating in another mode. In theory, this process is described at the Hamiltonian level by the non-linear contribution

$$\mathcal{H}_{XPM} = \chi a_1^\dagger a_1 a_2^\dagger a_2, \quad (1.9)$$

in which a_1 and a_2 are the two modes propagating through the medium, and χ represents the non-linear energy scale associated to the n_2 coefficient.

In the *self-phase modulation* regime, a field instead acquires a phase shift proportional to its own intensity, which at the quantum level can be described by the non-linear Hamiltonian contribution

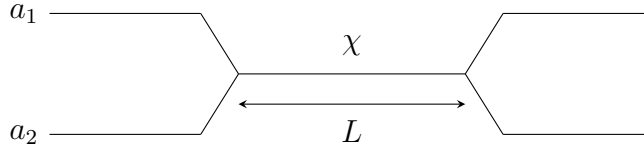
$$\mathcal{H}_{SPM} = \chi (a^\dagger a)^2. \quad (1.10)$$

While SPM can in principle generate non-classical states, such as Schrödinger cat states, both regimes share the same fundamental drawback: the non-linear effect is extremely weak at the single-photon level in most conventional materials known, to date.

For a non-linear crystal of length L , the XPM Hamiltonian (cross-Kerr Hamiltonian) in Eq. 1.9 generates the unitary transformation

$$K = e^{-\frac{i\chi L}{v_g} a_1^\dagger a_1 a_2^\dagger a_2}, \quad (1.11)$$

schematically represented in the following sketch:



In principle, a phase shift of π would be sufficient to implement a deterministic controlled-Z (CZ) gate between two photonic qubits, making the Kerr effect an attractive mechanism for scalable PQC, at least in theory. Let us briefly outline how this photonic two-qubit gate would work, on paper.

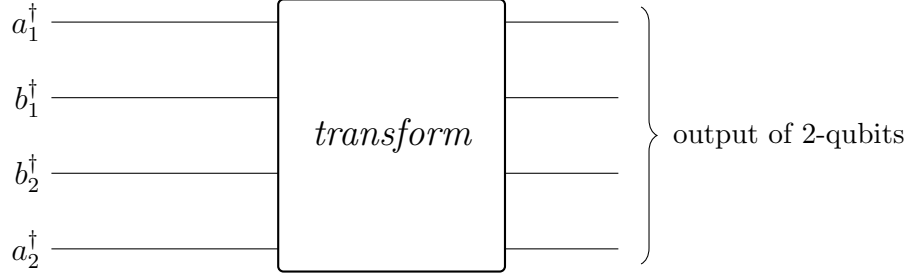
We assume a dual-rail qubit encoding with single-photon input states, for which the computational basis of the two qubits would transform, under the unitary evolution K , as:

$$K |00\rangle = |00\rangle, \quad K |01\rangle = |01\rangle, \quad K |10\rangle = |10\rangle, \quad K |11\rangle = e^{-i\chi L/v_g} |11\rangle.$$

Taking $\chi L/v_g = \pi$ and considering two dual-rail qubits implemented into four different optical modes, the relevant subspace is spanned by the four basis states (labelled on the occupancy of the 4 input channels) $|a_1, b_1, b_2, a_2\rangle$:

$$\{|1001\rangle, |0101\rangle, |1010\rangle, |0110\rangle\}.$$

Then, we assume that a Kerr medium is only present between the two middle optical modes, such that every state is left unchanged except for $|0110\rangle \xrightarrow{K} -|0110\rangle$. In this way, there is always a way to perform a global transformation on the input states, as schematically represented below, that would implement exactly the CZ gate.



Since, by the definition of the CZ gate (1.7), the CNOT matrix decomposes as

$$CNOT = (I \otimes H) CZ (I \otimes H), \quad (1.12)$$

we conclude that a CNOT gate can always be constructed from Kerr media, together with arbitrary single-qubit operations realized via beam splitters and phase shifters. Unfortunately, this route is currently out of reach, since no known Kerr medium is able to impart a π -phase shift at the single-photon level. In addition, in Refs. [3] and [4] it is argued that single-photon states are inevitably temporally limited wavepackets, which intrinsically constrains the achievable XPM effect:

- when the $\chi^{(3)}$ response time is shorter than the pulse duration, χ_{XPM} is negligible and cross-phase modulation imparts essentially no net phase shift;
- when the $\chi^{(3)}$ response time is longer than the pulse duration, the phase noise introduced by the non-instantaneous nature of the $\chi^{(3)}$ non-linearity strongly reduces the gate fidelity, making it impossible to simultaneously achieve a large (π) phase shift and high CZ fidelity.

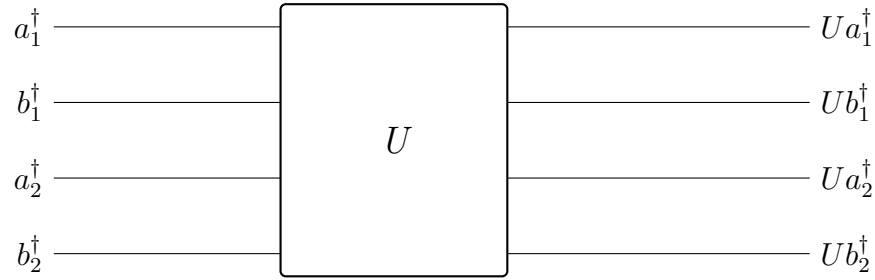
A further subtlety, shown explicitly in Ref. [3], is that even when a large phase shift appears in the Heisenberg picture, this does not translate into a correspondingly large phase shift in the Schrödinger picture: spontaneous emission populates initially unoccupied temporal modes, thus degrading the fidelity of the operation. Large phase shifts and high gate fidelity are therefore mutually incompatible requirements. Moreover, although the Kerr effect occurs in a wide range of materials, it remains extremely weak at the single-photon level, and the accumulated phase shift cannot be increased arbitrarily by extending the propagation length L , since most Kerr media are simultaneously affected by significant absorption or scattering losses that degrade the spatial mode of the propagating field.

These fundamental limitations, the incompatibility between large phase shifts and high gate fidelity, together with unavoidable absorption losses, have essentially ruled out classical non-linear optics as a viable route to scalable photonic quantum computation, so far. Nevertheless, as Knill, Laflamme, and Milburn showed in 2001 [5], the non-linearity required for two-qubit gates need not come from a material medium at all: it can be induced by measurement. We will outline this historically relevant paradigm in the following Section.

1.3.3 Linear optical quantum computing: the KLM paradigm

The so called LOQC paradigm rests on three physical resources: linear optical elements, single-photon sources, and photon-number-resolving detectors. At first sight, this seems insufficient for universal quantum computation, and indeed, it is impossible to implement two-qubit gates deterministically using linear optics alone, as we will hereby show. However, the historically remarkable insight from Knill, Laflamme, and Milburn [5] has been that this limitation can be circumvented: by allowing gate operations to be probabilistic and conditioned on measurement outcomes, photodetection itself provides an effective non-linearity that is sufficient to complete a universal gate set. In fact, this can be easily proved.

Let us consider a generic unitary transformation U acting on a two-qubit system encoded into four modes, respectively. Schematically, the qubits a and b and the generic transformation acting on them can be represented as follows



Now, we want to show that it is impossible to implement a deterministic entangling two-qubit gate by only using linear optical transformations.

Consider, for instance, the transformation

$$a_1^\dagger a_2^\dagger |\Omega\rangle \xrightarrow{U} \frac{1}{\sqrt{2}} (b_1^\dagger a_2^\dagger + a_1^\dagger b_2^\dagger) |\Omega\rangle ,$$

which maps a separable two-photon input to an entangled output state. However, a linear transformation U can only perform the following operations on the input fields:

$$\begin{aligned} a_1^\dagger &\xrightarrow{U} \alpha_1 a_1^\dagger + \beta_1 b_1^\dagger + \alpha_2 a_2^\dagger + \beta_2 b_2^\dagger \\ a_2^\dagger &\xrightarrow{U} \gamma_1 a_1^\dagger + \delta_1 b_1^\dagger + \gamma_2 a_2^\dagger + \delta_2 b_2^\dagger . \end{aligned}$$

The reason for that is rooted into the bosonic nature of photons and the linearity of the optical transformation: a linear optical network acts on creation operators via a fixed unitary matrix, M , mapping $a_i^\dagger \rightarrow \sum_j M_{ij} a_j^\dagger$. As a consequence, a product input state of the form $a_1^\dagger a_2^\dagger |\Omega\rangle$ is always mapped to another product of single-mode superpositions. In other words, linear optics cannot generate entanglement from a separable input, there is no combination of the eight coefficients such that the desired transformation happens. In fact:

$$(a_1^\dagger a_2^\dagger)_{in} \xrightarrow{U} \left(\sum_{i=1}^2 (\alpha_i a_i^\dagger + \beta_i b_i^\dagger) \right)_{out} \times \left(\sum_{i=1}^2 (\gamma_i a_i^\dagger + \delta_i b_i^\dagger) \right)_{out} \neq \frac{1}{\sqrt{2}} (b_1^\dagger a_2^\dagger + a_1^\dagger b_2^\dagger) .$$

The remarkable insight of Knill, Laflamme, and Milburn [5] is that this limitation can be circumvented: by allowing gate operations to be probabilistic and conditioned on

measurement outcomes, photodetection itself provides an effective non-linearity that is sufficient to complete a universal gate set.

Essentially, the KLM scheme (named so after the authors' initials) relies on two key ingredients. The first is the bosonic nature of photons, which causes two indistinguishable photons to bunch together at a beam splitter, i.e., the celebrated Hong-Ou-Mandel (HOM) effect. The second is the consequence of projective measurements on ancilla modes, which collapses the photonic state in a way that mimics a non-linear interaction.

We hereby summarize the main ingredients leading to the definition of probabilistic two-qubit gates within LOQC. We start by introducing elementary linear optical operations acting on single-photon Fock states, such as beam splitters, phase shifters, and waveguides. Altogether, these elements allow for an arbitrary manipulation of photonic states. Then, we will proceed by defining the HOM effect, and its relevance to the implementation of probabilistic two-qubit gates.

Waveguides. Optical waveguides, typically realized in high index insulating or semi-conducting materials, satisfy the need for a deterministic guiding quantum states of light. Normally, signals are injected into waveguides in the form of single photons, or photon wave packets, thus ensuring that they interact with other optical elements in a controlled and reliable fashion. They come in a variety of forms, including optical fibers and integrated photonic circuits in silicon (Si), silicon nitride (SiN), glass, and more recently Lithium niobate. In the context of LOQC, on-chip integrated waveguides are particularly advantageous, as they enable precise control over photon routing and propagation. These structures can be engineered to support specific propagation modes, such as single-mode or multi-mode operations, and can be combined with other optical components to implement quantum gates and more complex quantum circuits. Henceforth, for all the purposes of the present thesis we will assume that single-mode low loss propagating modes can be engineered in photonic integrated circuits, without further specifying the material platform or the details of the electromagnetic engineering of the various optical components.

Beam Splitters. Beam splitters (or directional couplers) are fundamental linear optical elements that distribute photons across two distinct output paths. In quantum optics, a beam splitter is described by a unitary transformation acting on the quantum state of the incoming photon. A prototypical example is the 50/50 (or *balanced*) beam splitter, in which an incoming photon is redirected with equal probability into either of the two output paths. From a theoretical point of view, a 50/50 beam splitter acts on the annihilation operators a_1 and a_2 of the two input modes as follows:

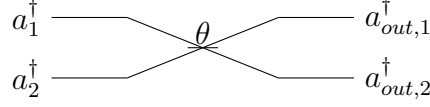
$$\begin{aligned} a_{out,1}^\dagger &= \frac{1}{\sqrt{2}}(a_1^\dagger + a_2^\dagger) \\ a_{out,2}^\dagger &= \frac{1}{\sqrt{2}}(a_1^\dagger - a_2^\dagger) \end{aligned} \tag{1.13}$$

where a_1 and a_2 are the annihilation operators for the input modes, and $a_{out,1}$ and $a_{out,2}$ are the annihilation operators for the two output modes, respectively.

In the general case, the beam splitter can be parametrized by a single real number, i.e, an

index θ representing the mixing angle of the directional coupler, which is in turn related to the reflectivity (transmissivity) of the optical element as $R = \sin^2 \theta$ ($T = \cos^2 \theta$), with $R + T = 1$. Hence, in the general case the beam splitter transformation (and the corresponding optical circuit element) can be given as:

$$a_{out,1}^\dagger = \cos(\theta)a_1^\dagger + i \sin(\theta)a_2^\dagger \quad a_{out,2}^\dagger = i \sin(\theta)a_1^\dagger - \cos(\theta)a_2^\dagger$$



An equivalent (and very useful) representation of this problem is related to the definition of a Hamiltonian, the *beam-splitter Hamiltonian*, which can be written (up to a phase factor) as:

$$H_{BS} = i\theta \left(a_1^\dagger a_2 + a_1 a_2^\dagger \right). \quad (1.14)$$

In fact, the Hamiltonian in Eq. (1.14) simply produces the following transformation between input and output modes:

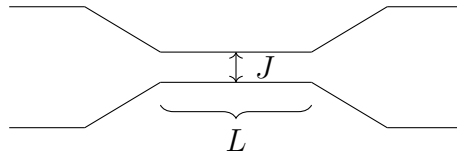
$$\begin{bmatrix} a_{out,1}^\dagger \\ a_{out,2}^\dagger \end{bmatrix} = \begin{bmatrix} \cos \theta & i \sin \theta \\ i \sin \theta & \cos \theta \end{bmatrix} \begin{bmatrix} a_1^\dagger \\ a_2^\dagger \end{bmatrix} = (\cos(\theta)\mathbb{I} + i \sin(\theta)\sigma_x) \begin{bmatrix} a_1^\dagger \\ a_2^\dagger \end{bmatrix} = R_x(-2\theta). \quad (1.15)$$

Essentially, in quantum computing terminology, this would correspond to a rotation around the axis x by an angle θ . By tuning θ , and arbitrary rotation of the single-qubit state around x can be implemented.

In the field of integrated Photonics, the very same transformation is straightforwardly interpreted as the typical tunneling-type Hamiltonian arising, e.g., from two coupled waveguides in an integrated circuit:

$$H = J(a_1^\dagger a_2 + a_1 a_2^\dagger),$$

in which J represents the tunnel energy, and the rotation angle θ is intrinsically linked to the propagation time within the waveguides, such as $\theta = Jt/\hbar = JL/(\hbar v_g)$, where L is the length of the tunnel-coupling region, and v_g is the group velocity of the single-photon wave packet as schematically shown below.

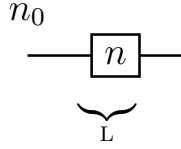


Beam splitters are useful in LOQC for tasks such as photon interference and entanglement generation. Crucially, the interference of two indistinguishable photons at a beam splitter is not merely a classical wave phenomenon: it reflects the bosonic nature of photons and gives rise to strongly non-classical correlations. As we will discuss in the context of the KLM scheme, this quantum interference — the already mentioned HOM effect — can be exploited, together with projective measurements, to simulate an effective photon-photon non-linearity without requiring any non-linear medium.

Phase Shifters. Phase shifters are optical elements designed to introduce a controlled delay in the phase of propagating light beams. In quantum information processing, phase shifts play a vital role in manipulating the relative phases of quantum states, thereby modifying interference patterns and enabling the generation of entanglement. In the LOQC framework, phase shifters typically consist of optical materials, such as high-index semiconducting materials, or glass, whose refractive index can be tuned to impart a desired phase shift to a passing photon. More precisely, a phase shifter applies a shift $\Delta\phi$ to the photon mode, transforming the annihilation operator according to:

$$a_{out} = e^{i\Delta\phi} a_{in} ,$$

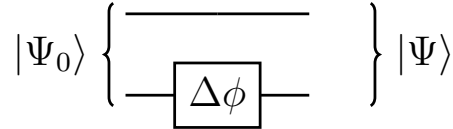
in which a_{in} and a_{out} denote the annihilation operators of the input and output modes, respectively. A single propagating photon experiences a spatial phase e^{ikL} , with the wave vector $k = n\omega/c$, in which n is the modulated refractive index in the region of length L , while n_0 is the refractive index of the background medium. In this way, the overall phase shift is given as: $\Delta\phi = (n - n_0)\omega L/c$.



At the Hamiltonian level, a phase shifter just acts as a normal time evolution operator, with $H_P = \Delta\phi a^\dagger a \rightarrow P = e^{i\Delta\phi a^\dagger a}$ such that:

$$\begin{aligned} P|0\rangle_a &= |0\rangle_a & P|1\rangle_a &= e^{i\Delta\phi}|1\rangle_a \\ P &= e^{i\frac{\Delta\phi}{2}} \begin{bmatrix} e^{-i\frac{\Delta\phi}{2}} & 0 \\ 0 & e^{+i\frac{\Delta\phi}{2}} \end{bmatrix} = R_z(\delta\phi) . \end{aligned} \quad (1.16)$$

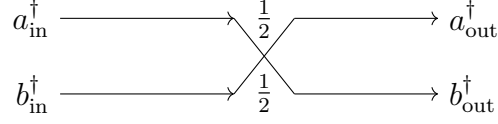
Hence, in quantum computing terminology, a phase shifter produces a rotation of the qubit state around the z axis. In particular, by assuming dual rail qubit encoding, we have the following circuit description of the phase shifter:



$$|\Psi_0\rangle = \alpha|10\rangle + \beta|01\rangle \quad |\Psi\rangle = \alpha e^{-i\omega t}|10\rangle + \beta e^{-i\omega t + i\Delta\phi}|01\rangle .$$

In fact, by assuming that a controlled phase shift can be imparted, we can say that an arbitrary rotation around z can be implemented in PQC. More generally, phase shifters are employed to implement quantum gates such as the phase gate, and to govern the interference between different photon paths in interferometric configurations, which are central to LOQC protocols such as Boson Sampling.

The Hong-Ou-Mandel effect. The so called Hong-Ou-Mandel (HOM) effect is a striking manifestation of the bosonic nature of light quanta [28, 29]. Here we give a brief of its main features. We consider a balanced directional coupler (i.e., a 50/50 beam-splitter), schematically represented as the circuit below



which can be parametrized with the following transfer operation between input and output field operators:

$$\begin{pmatrix} a_{\text{out}}^\dagger \\ b_{\text{out}}^\dagger \end{pmatrix} = \begin{pmatrix} \cos \frac{\pi}{4} & i \sin \frac{\pi}{4} \\ i \sin \frac{\pi}{4} & \cos \frac{\pi}{4} \end{pmatrix} \begin{pmatrix} a_{\text{in}}^\dagger \\ b_{\text{in}}^\dagger \end{pmatrix} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & i \\ i & 1 \end{pmatrix} \begin{pmatrix} a_{\text{in}}^\dagger \\ b_{\text{in}}^\dagger \end{pmatrix}$$

If we now assume the quantum state $|11\rangle$ as the input, this is then transformed as:

$$a_{\text{in}}^\dagger b_{\text{in}}^\dagger \rightarrow \frac{1}{2}(a_{\text{in}}^\dagger + i b_{\text{in}}^\dagger)(i a_{\text{in}}^\dagger + b_{\text{in}}^\dagger) = \frac{i}{2}(a_{\text{in}}^{2\dagger} + b_{\text{in}}^{2\dagger}) + \frac{1}{2}(a_{\text{in}}^\dagger b_{\text{in}}^\dagger - a_{\text{in}}^\dagger b_{\text{in}}^\dagger) = \frac{i}{2}(a_{\text{in}}^{2\dagger} + b_{\text{in}}^{2\dagger})$$

So, ultimately, the input state is transformed as $|11\rangle \rightarrow \frac{1}{2}(|20\rangle + |02\rangle)$, up to a global and irrelevant phase.

The amplitude of the $|11\rangle$ state vanishes exactly at the output due to destructive interference between the two indistinguishable bosonic paths [29]. Essentially, this means that two indistinguishable photons entering a 50/50 beam splitter will always exit together in the same output mode — a coincidence click (one photon in each output) has zero probability. This “bunching” behavior is a purely quantum effect with no classical analogue, and it is the key physical mechanism underlying the KLM scheme.

The non-linear sign gate. Another key ingredient in this construction is the so-called *non-linear sign gate* (NS gate), i.e., a single-mode operation that acts as the identity on $|0\rangle$ and $|1\rangle$ with respect to photons propagating in the system, while mapping $|2\rangle$ to $-|2\rangle$, as follows.

$$\text{NS} : \quad |0\rangle \mapsto |0\rangle, \quad |1\rangle \mapsto |1\rangle, \quad |2\rangle \mapsto -|2\rangle$$

This gate cannot be implemented deterministically with linear optics alone, but Knill, Laflamme and Milburn showed that it can be realized probabilistically using 2 ancilla photons, a network of beam splitters, and projective measurements on the ancilla modes, succeeding with probability $p_{\text{NS}} = 1/4$.

A scheme of the linear optical network allowing to implement the NS gate is represented in Fig. 1.2. Essentially, the quantum circuit acts on three modes: mode 1 carries the input state, generically given as $|\psi\rangle = \alpha_0 |0\rangle + \alpha_1 |1\rangle + \alpha_2 |2\rangle$, while modes 2 and 3 are ancillas prepared in the Fock states $|1\rangle$ and $|0\rangle$, respectively. The network consists of three directional couplers (i.e., unbalanced beam splitters), respectively parametrised by angles (θ_1, ϕ_1) , (θ_2, ϕ_2) and (θ_3, ϕ_3) , and a phase shifter ϕ_4 acting on mode 1. The output is accepted only upon a successful post-selection event, in which the single-photon

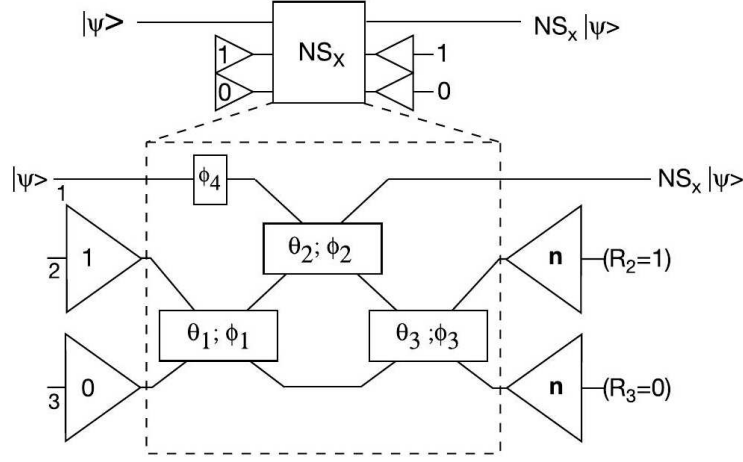
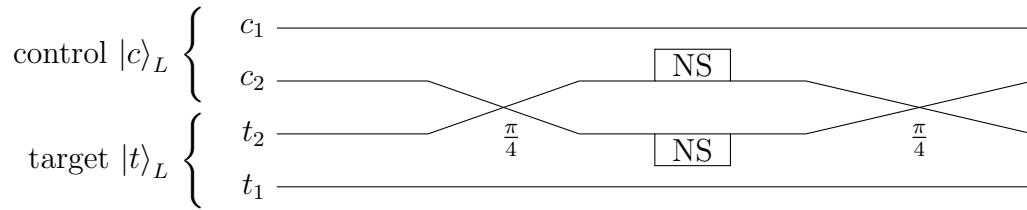


Figure 1.2: Scheme of the circuit allowing to implement a non-linear phase shifts on a given input state in one mode. Schematic figure adapted from Ref. [5].

photon counters register exactly one photon in mode 2 ($R_2 = 1$) and zero photons in mode 3 ($R_3 = 0$); the combination of these events is shown to occur with probability $1/4$. When successful, the gate realises the transformation $\alpha_0 |0\rangle + \alpha_1 |1\rangle + \alpha_2 |2\rangle \rightarrow \alpha_0 |0\rangle + \alpha_1 |1\rangle - \alpha_2 |2\rangle$, i.e., a conditional sign flip on the two-photon component. For the specific case $x = -1$, the required angles are $\theta_1 = 22.5^\circ$, $\phi_1 = 0^\circ$, $\theta_2 = 65.5302^\circ$, $\phi_2 = 0^\circ$, $\theta_3 = -22.5^\circ$, $\phi_3 = 0^\circ$ and $\phi_4 = 180^\circ$. The dashed box defines the abbreviated symbol NS_x used to represent this gate from here on. Here we take the NS gate as a primitive and use it to construct a probabilistic CZ gate.

CZ gate from NS gates. A CZ gate on two dual-rail encoded qubits, labelling the four modes as c_1, c_2 (control) and t_1, t_2 (target), can be constructed as follows [5], according to the scheme represented below.



The input state is assumed to be the generic:

$$|c\rangle_L \otimes |t\rangle_L$$

in which c and t stand for control and target (qubits), respectively, and the logical states of the two qubits are defined, given the vacuum state $|\Omega\rangle$, as follows:

$$\begin{cases} |0\rangle_c = \hat{c}_1^\dagger |\Omega\rangle \\ |1\rangle_c = \hat{c}_2^\dagger |\Omega\rangle \end{cases} \quad \begin{cases} |0\rangle_t = \hat{t}_1^\dagger |\Omega\rangle \\ |1\rangle_t = \hat{t}_2^\dagger |\Omega\rangle \end{cases}$$

The circuit above applies a $\pi/4$ beam splitter between modes c_2 and t_2 , followed by a NS gate on each of those modes, followed by a second $\pi/4$ beam splitter. The modes

c_1 and t_1 pass through unchanged.

Let us now consider a general logical input, the sequences of transformations work as follows. The general input state is:

$$\begin{aligned} |\psi_{in}\rangle &= |c\rangle_L \otimes |t\rangle_L \rightarrow (\alpha|10\rangle_c + \beta|01\rangle_c) \otimes (\gamma|01\rangle_t + \delta|10\rangle_t) = \\ &= \alpha\gamma \underbrace{|1001\rangle}_{|00\rangle_L} + \alpha\delta \underbrace{|1010\rangle}_{|01\rangle_L} + \beta\gamma \underbrace{|0101\rangle}_{|10\rangle_L} + \beta\delta \underbrace{|0110\rangle}_{|11\rangle_L}. \end{aligned}$$

Here, the terms $|1001\rangle$, $|1010\rangle$, $|0101\rangle$ contain at most one photon in the $\{c_2, t_2\}$ subspace, so the beam splitter simply redistributes that photon and leaves the logical content unchanged (up to single-photon amplitudes that will cancel in the second beam splitter). Only the $|0110\rangle$ term, which contributes one photon in each of c_2 and t_2 , is affected non-trivially: as a consequence of the HOM effect, the two photons bunch into a two-photon Fock state. After that, the NS gate maps $|2\rangle \mapsto -|2\rangle$. Finally, the second beam splitter is the inverse of the first: it maps the two-photon bunched state back to the logical $|11\rangle$ term, but now with a relative minus sign introduced by the NS gates.

$$\begin{aligned} |\psi_{in}\rangle &\xrightarrow{\pi/4\text{-BS}} \alpha\gamma|00\rangle_L + \alpha\delta|01\rangle_L + \beta\gamma|10\rangle_L + \frac{\beta\delta}{2}(|0200\rangle + |0020\rangle) \xrightarrow{\text{NS}} \\ &\xrightarrow{\text{NS}} \alpha\gamma|00\rangle_L + \alpha\delta|01\rangle_L + \beta\gamma|10\rangle_L - \frac{\beta\delta}{2}(|0200\rangle + |0020\rangle) \xrightarrow{\pi/4\text{-BS}} \\ &\xrightarrow{\pi/4\text{-BS}} \alpha\gamma|00\rangle_L + \alpha\delta|01\rangle_L + \beta\gamma|10\rangle_L - \beta\delta|11\rangle_L \end{aligned}$$

Hence, as a final outcome we obtain the well known result of a CZ gate, i.e., the one represented by the matrix transformation

$$|\Psi\rangle_{out} = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & -1 \end{bmatrix} |\Psi_{in}\rangle.$$

Notice that the overall success probability is $p_{CZ} = (1/4)^2 = 1/16$, since the two NS gates act in parallel, and each succeeds independently with probability $1/4$. Since we have already mentioned that the CNOT gate can be straightforwardly obtained as $CNOT = (I \otimes H) CZ (I \otimes H)$, and Hadamard gates are deterministically implementable with linear optics, the CNOT gate inherits the same overall success probability. This probabilistic character is a fundamental feature of all two-qubit gates implementations in LOQC, and overcoming it (through gate teleportation, resource state engineering, or measurement-based paradigms), which is the central aspect faced in the present thesis.

1.3.4 State of the art in Photonic Quantum Computing

Controlling individual quantum systems (such as photons, atoms, or spins) lies at the heart of quantum technologies. A key enabling technology comprises efficient tools to generate quantum states of light. In this respect, a widely used category of integrated quantum sources relies on non-linear optical phenomena, including spontaneous parametric down-conversion (SPDC). Key strengths of this approach are the capacity to multiplex photon states at high speed and the fact that the emitted photons

are heralded and share the same wavelength. On the downside, these sources produce photons at relatively low rates and operate in an inherently probabilistic fashion. In contrast, on-demand single-photon sources (SPSs) achieve higher generation efficiency, and certain solid-state SPSs, such as defects in wide-bandgap materials or semiconductor quantum dots (QDs), can host a spin degree of freedom, which plays a key role in entanglement distribution and quantum network implementation.

Progress in nanofabrication over recent decades has made it possible to integrate quantum light sources (both SPDC sources and single QDs) with on-chip nanophotonic elements, such as optical resonators and waveguides. Nevertheless, reconfiguring individual photonic components on a chip remains essential for executing different quantum protocols without requiring hardware changes. Fast source modulation (at GHz rates) and the realization of multiple indistinguishable sources continue to pose significant challenges.

Thus far, the only photonic components showing reliable and efficient modulation are phase shifters, implemented in materials such as Lithium Niobate, silicon nitride, or silicon-on-insulator (SOI) platforms. Fully programmable, integrated quantum photonic systems could open a viable path toward large-scale quantum processors characterized by high bandwidth, low energy consumption, and minimal system losses. Integrated quantum photonics is indeed a rapidly growing field that brings together quantum light sources, non-linear materials, photonic resonators, and single-photon detectors onto a single chip. One of the most important directions in this area is the push toward fully programmable devices — systems that can be quickly reconfigured without any hardware modifications.

In this Section, we will briefly mention the ongoing research directions that are actively pursuing the development of tunable elements (such as phase shifters and quantum light sources), and new materials (such as van der Waals crystals), which are opening up fresh possibilities for on-chip control. To do so, we will mainly rely on Refs. [25, 30].

Single-photon Sources.

A single-photon state can be described in terms of a multimode wavepacket as

$$|1; f\rangle = \sum_{m=1}^{\infty} f_m \hat{a}_m^\dagger |0\rangle,$$

where the amplitudes f_m satisfy $\sum_m |f_m|^2 = 1$ and characterize the temporal and spectral profile of the emitted photon. A key diagnostic for a single-photon source is the suppression of two-photon coincidences, quantified by the second-order correlation function $g^{(2)}(\tau)$: a good source must exhibit $g^{(2)}(0) \approx 0$, reflecting photon antibunching. As already mentioned, the most widely used class of integrated quantum sources relies on SPDC, in which a pump laser drives a non-linear crystal to produce correlated photon pairs. The output state can then be written as

$$|\Psi_{\text{PDC}}\rangle = \sqrt{1 - |\lambda|^2} \sum_{n=0}^{\infty} \lambda^n |n, n\rangle,$$

in which $|\lambda|^2$ controls the pair generation probability. Key strengths of this approach are the capacity to multiplex photon states at high speed, and the fact that emitted

photons are heralded and share the same wavelength. However, these sources operate in an inherently probabilistic fashion, and intrinsically suffer from multi-photon contamination at higher pump powers.

Alternatively, on-demand *single-photon sources* (SPSs) achieve higher generation efficiency. In this respect, a second class of deterministic sources is based on *cavity-QED Raman schemes*, in which a single emitter is driven by a pulsed laser to emit exactly one photon into a cavity mode via a stimulated Raman process. The emission time is fully controlled by the pump pulse, enabling precise wavepacket shaping. Unlike SPDC, Raman-based sources produce photons with a Gaussian temporal profile, which is significantly more robust to mode mismatching than the Lorentzian profile characteristic of spontaneous emission([31]). Certain solid-state SPSs (such as defects in wide-bandgap materials or semiconductor QDs) can additionally host a spin degree of freedom, which plays a key role in entanglement distribution and quantum network implementation. Since any two solid-state SPSs are inherently non-identical, significant efforts are directed toward achieving rapid in-situ frequency tuning via strain or applied electric fields [32, 33].

A particularly promising frontier is represented by SPSs derived from *van der Waals* (vdW) *crystals*, which offer new degrees of freedom unavailable in conventional 3D crystals, such as deterministic assembly and twist-angle control. The emission wavelength can be adjusted by varying the twist angle or by choosing a different colour centre within the same host material. Similarly, vdW crystals can be employed for non-linear applications, including the generation of SPDC sources (with rhombohedral boron nitride (rBN) and transition metal dichalcogenides representing notable examples [34]) and allow the engineering of periodically poled microscale crystals for efficient frequency conversion, orders of magnitude thinner than conventional non-linear crystals([35]).

For any LOQC scheme, multiple sources must produce *identical* photons, since indistinguishability is required for high-visibility quantum interference. This characteristic is tested via the Hong-Ou-Mandel (HOM) effect: only perfectly indistinguishable photons produce complete suppression of coincidence events at a 50/50 beam splitter, making HOM visibility a standard figure of merit for single-photon source quality.

Structured light.

Structured light is defined as spatial modes varying in amplitude, phase, and polarization, including spatially multimodal quantum states, quantum vortices and vector beams. It can be highly beneficial for a variety of integrated quantum systems. It has the potential to increase information storage capacity through high-dimensional quantum states (qudits), to provide enhanced coupling to spin or valley degrees of freedom (particularly within 2D materials), and to facilitate efficient routing and multiplexing, which is highly sought after in quantum circuitry ([36, 37]). The qudit implementation, in particular, is not exploited in the present thesis, but may become a route to possible extensions of this work.

The generation of quantum structured light is still in its early stages: the main challenges lie in fabricating the nanostructures needed to produce structured light on chip, in the complexity of the detection circuitry required, and in maintaining phase stability and mode purity [38]. Tunability during generation remains an open challenge. Greater flexibility for dynamically tuning quantum states of structured light is offered when pho-

tons are generated through SPDC and spontaneous four-wave mixing (SFWM), as these processes allow the generation of entangled multi-photon states in CMOS-compatible photonic circuitry [39].

Dynamic modulation of on-chip components.

Dynamic manipulation of on-chip components is a critical building block for programmable quantum photonics. Frequency shifters and beam splitters with low on-chip losses have historically been difficult to realize, as traditional acousto- or electro-optic approaches tend to be bulky, non-bidirectional, or lossy. Phase shifters play a particularly central role: when integrated into tunable couplers and splitters, they enable both phase and amplitude modulation. Key requirements include high switching speed, minimal optical loss, low power consumption, and compatibility with cryogenic temperatures — the latter being essential for the integration of single-photon sources and detectors.

Several technologies have been explored in this respect, each with distinct trade-offs: thermo-optic shifters are simple and CMOS-compatible but slow and power-hungry; free-carrier modulators improve speed but introduce optical losses; electro-optic designs based on lithium niobate or barium titanate offer high speed but come with integration complexity; MEMS-based approaches provide compact, low-power solutions; and phase-change materials enable non-volatile reconfiguration with zero static power [40, 41, 42, 43].

The *thermo-optic phase shifter* controls the phase by locally heating it the waveguide where the photons are propagating. The working principle relies on the thermo-optic effect ([44]): in many materials, such as silicon or silicon nitride, the refractive index depends on the material temperature. When a small resistive heater (typically a metallic wire on top of the waveguide) is driven by an electrical current, it raises the local temperature, which in turn changes the refractive index of the waveguide material. Since the phase accumulated by a light wave travelling through a medium depends on the refractive index (see discussion above), this temperature change translates directly into a controllable phase shift. By tuning the applied voltage or current, one can continuously and precisely adjust the phase of the optical mode. Future phase-shifter solutions will likely combine multiple mechanisms and heterogeneous material integration to meet the full set of requirements.

Single-photon detectors.

Single-photon detection provides the projective measurements that underlie both gate operations and readout in LOQC. Two broad classes of detectors are relevant: *number-resolving detectors*, which can discriminate between different photon numbers, and *bucket detectors*, which produce only a binary output without resolving photon number. Real detectors are subject to two main error mechanisms: *photon loss*, where fewer photons are registered than were present, and *dark counts*, where a detection event occurs in the absence of any input photon. The detector efficiency $\eta \in [0, 1]$ is defined as the probability that a single-photon input produces a detection event.

The most common detectors in LOQC experiments are *avalanche photodiodes* (APDs), which operate as bucket detectors. An incoming photon triggers an electron avalanche

in the active semiconductor region, producing a measurable current. APDs achieve typical efficiencies of around 85% at 660nm, with dark count rates as low as ~ 25 Hz at cryogenic temperatures, but cannot resolve photon number due to their dead time of a few nanoseconds.

Transition-edge sensors (TES) offer photon-number resolution via superconducting calorimetry: incoming photons raise the temperature of an absorber, which is measured with high precision. TES detectors operate below 100 mK and achieve efficiencies exceeding 88% with negligible dark counts, at the cost of a slow repetition rate of the order of 10 kHz imposed by the thermal reset time.

Finally, *detector cascading* provides an approximate number-resolving capability by splitting an incoming mode across N output ports, each monitored by a bucket detector. A time-domain variant known as *time multiplexing* achieves the same effect using fiber delay loops, though the required fiber length grows exponentially with the desired photon-number resolution [45, 46].

Programmable quantum photonic processing.

An early and significant example of a reconfigurable photonic platform for quantum computing is represented by the work from J. Carolan et al. (2015) (see Ref. [47]). Here, the authors demonstrated a single reprogrammable optical circuit built on a photonic chip, capable of implementing any linear optical protocol up to the size of the circuit itself. The device consisted of six optical modes connected through a cascade of 15 Mach-Zehnder interferometers, each controlled by thermo-optic phase shifters, allowing arbitrary unitary transformations to be programmed electrically. Rather than fabricating a dedicated chip for each specific quantum protocol, a single universal device could be reconfigured to run different experiments — from heralded quantum logic gates to boson sampling complex Hadamard operations — achieving an average fidelity of 0.999 ± 0.001 over 100 random unitary matrices. This work established an important proof of concept: universality and reconfigurability could coexist in an integrated photonic platform.

Building on this foundational work, more recent developments have pushed photonic quantum computing closer to practical usability. A notable example is the Ascella platform developed by the private French company Quandela [24]. Ascella is a cloud-accessible quantum computing machine based on six single photons generated by a high-efficiency semiconductor quantum dot source, which offers a significant advantage over the SPDC sources used in earlier experiments, as it produces photons on demand with higher purity and indistinguishability (generated by a single electrically controlled semiconductor QD). The photons are routed through a 12-mode reconfigurable universal interferometer, and fully controlled by 126 thermo-optic phase shifters. A distinctive feature of this platform is the use of a machine-learned transpilation process to compensate for hardware imperfections — an approach that reflects the growing role of artificial intelligence in the optimization of photonic chips — achieving an average matrix implementation fidelity of $99.7\% \pm 0.08\%$. The material used for the photonic chip is SiN. The platform supports both gate-based and photon-native computation paradigms, and benchmarks record fidelities for one, two and three qubit gates. Taken together, these two works illustrate the trajectory of integrated photonic quantum computing: from an early laboratory demonstration of universality to a multi-purpose, remotely acces-

sible device capable of running a variety of quantum algorithms with state-of-the-art performance.

As a last comment, it is worth noting that both platforms operate entirely in the linear optical regime, without exploiting any non-linear interaction between photons. This is a deliberate and practically motivated choice: strong optical non-linearities are extremely difficult to engineer at the single-photon level. Instead, both works rely on the Knill-Laflamme-Milburn (KLM) scheme ([5]), in which the effective non-linearity needed to implement two-qubit gates is induced probabilistically through measurements on ancilla photons: when these are detected in the correct output modes, the gate is confirmed to have succeeded — a process known as heralding. The price to pay is that such gates succeed only with a certain probability: for instance, the CNOT gate in Ascella succeeds with probability $1/9$, and the Toffoli gate with approximately $1/57$. This probabilistic nature represents one of the main limitations of the linear-optical approach. Scaling it to a fault-tolerant regime will likely require moving towards measurement-based paradigms that rely on pre-generated multipartite entangled states — known as cluster states — as a universal resource for computation (highlighted in Sec. 1.5 of this thesis), rather than on probabilistic gate sequences alone. Or, as it will be discussed in the remainder of this thesis, we could aim for exploiting the small optical non-linearities in Kerr-type materials to achieve a higher fidelity computational paradigm still based on probabilistic gating.

Satisfying DiVincenzo Criteria in PQC.

With all the relevant information briefly outlined above, concerning the main aspects of the current status of PQC, we can now check the satisfaction of the DiVincenzo criteria for this quantum technological platform as follows:

- 1. Well-defined qubits:** Photonic qubits can be encoded using various degrees of freedom such as polarization, path, dual-rail, time-bin, or frequency. In dual-rail encoding, a single qubit is represented by a single photon distributed across two distinct spatial modes, enhancing robustness and compatibility with linear optical gates.
- 2. Qubit initialization:** Single-photon sources combined with a spatial demultiplexer can generate well-defined path-encoded qubits. Two main approaches are available: heralded single photons from SPDC sources, and on-demand single photons from deterministic sources such as semiconductor quantum dots or cavity-QED Raman schemes.
- 3. Long coherence times:** Photons interact weakly with their environment, giving them naturally excellent coherence properties. In free space or optical fibers, photons can maintain coherence over long distances, though propagation losses must be managed carefully.
- 4. Universal gate set:** Linear optical elements such as beam splitters and tunable phase shifters can implement arbitrary single-qubit operations. Two-qubit gates can be realized through quantum interference and ancillary photon measurements, as in the

KLM protocol — though they are inherently probabilistic, as discussed in detail in the following sections.

5. Qubit-specific measurement: Photonic qubits can be measured using several detector technologies. Avalanche photodiodes (APDs) are the most common in practice, operating as bucket detectors with efficiencies around 85%. Transition-edge sensors (TES) offer photon-number resolution with efficiencies exceeding 88%, while superconducting nanowire single-photon detectors (SNSPDs) combine high efficiency with low dark counts and fast timing response [48].

6. Interconversion capability: Photons naturally serve as flying qubits, making photonic systems particularly well-suited for interfacing with quantum communication networks and for distributing entanglement over long distances.

7. Faithful transmission: Photons can be transmitted through optical fibers or free space with relatively low decoherence. For long-distance communication, propagation losses must be managed through quantum repeaters or entanglement purification protocols.

In summary, we have outlined the main conceptual aspects as well as practical building blocks behind PQC, from the encoding of qubits in optical modes to the manipulation of photonic states via linear optical elements. We have seen that, while photons are ideal carriers of quantum information due to their low decoherence, the absence of natural photon-photon interactions that are sufficiently strong to impart a π -phase shift to a multi-photon wave-packet as compared to single photon ones poses a fundamental challenge for the implementation of deterministic two-qubit gates in PQC. The KLM scheme offers an historically elegant solution by exploiting quantum interference and projective measurements to simulate an effective non-linearity, at the cost of probabilistic gate success. Overcoming this limitation — either through improved linear-optical schemes, measurement-based paradigms, or the engineering of genuine optical non-linearities — remains one of the central open problems in the field.

1.4 Alternative quantum computing platforms

Quantum computing relies on the ability to physically realize and control two-level quantum systems, which can be used to encode information into suitable qubits. Several hardware platforms have been developed to this end, each exploiting different physical systems and materials to encode and manipulate quantum information. Beyond the photonic approaches outlined in the previous Section, the most experimentally mature platforms include trapped ions, trapped neutral atoms, superconducting circuits, nitrogen-vacancy (NV) centers in diamond, and semiconductor spin qubits [49]. Despite sharing the same underlying computational framework, these platforms substantially differ between each other, either in their physical implementation or in control mechanisms and scalability prospects. For completeness, we hereby report a brief summary of the most relevant platforms currently explored in international research laboratories.

In *trapped ions* (see, e.g., Refs. [50, 51]) and *trapped neutral atoms* (see Ref. [52]) platforms, quantum information is encoded in the internal electronic energy levels of single atoms, either ionized or neutral species (typically, two stable hyperfine ground states selected for their long intrinsic lifetimes). In trapped ion systems, the charged nature of the ions allows effective confinement via electric fields, and manipulation is achieved through microwave fields or external lasers; in neutral atoms, confinement relies on the induced dipole moment created by optical potentials, such as optical tweezers or lattices, and two-qubit interactions are typically mediated via highly excited Rydberg states.

Superconducting circuits (see Ref. [53]) follow a radically different approach: qubits are not natural atoms but macroscopic electrical circuits incorporating non-linear inductive elements, such as Josephson junctions. The two computational levels correspond to the two lowest energy eigenstates of these engineered “artificial atoms”, and manipulation is carried out via microwave pulses. This approach allows considerable flexibility in tailoring the energy level structure.

NV centers in diamond (see, e.g., Ref. [54]) exploit single point defects in the crystal lattice instead, in which a nitrogen atom replaces a carbon atom adjacent to a vacancy in the crystal lattice. The qubit is encoded in the electron spin of the NV center ground state, controlled via microwave fields, while nearby nuclear spins — such as those of the nitrogen atom itself (^{14}N or ^{15}N) or naturally occurring carbon-13 within the diamond lattice, can serve as auxiliary qubits or long-lived quantum memories.

Finally, *semiconductor spin qubits* (see, e.g., Ref. [55]) encode information in the spin of charge carriers confined in nanostructures: typically electrons or holes trapped in quantum dots, or alternatively the nuclear or electron spin of donor atoms such as phosphorus implanted in a silicon host. Although they share with NV centers the use of spin in a solid-state environment, they differ significantly in the host material and in the specific nature of the spin carrier.

All of these platforms represent promising candidates for the realization of fault-tolerant quantum computation, as demonstrated by the remarkable experimental progress documented over the past two decades. The stability of a physical qubit is fundamentally characterized by two timescales: T_1 , the energy relaxation time, which quantifies how long a qubit retains its energy before incoherently exchanging it with the environment, and T_2 , the dephasing time, which measures how long the relative phase between the two computational states is preserved in a superposition. Both have improved by orders of magnitude across all platforms, with superconducting circuits doubling their T_1 every approximately one year and trapped ions doubling their T_2 every 1.9 years. Entanglement error rates have also consistently fallen, with semiconductor spins halving their error rate every 1.2 years and ion traps achieving entanglement errors as low as 0.03%. In terms of scalability, neutral atom platforms have reached arrays of up to 6100 qubits, while superconducting circuits and trapped ions have demonstrated systems of 127 and 100 qubits respectively. Recent milestones such as the demonstration of surface code operation below the fault-tolerance threshold on a 101-qubit distance-7 superconducting processor further confirm that these platforms are steadily converging toward the requirements for practical quantum computation. [56]

While significant challenges remain, the trajectory of progress across all platforms gives strong reason for optimism.

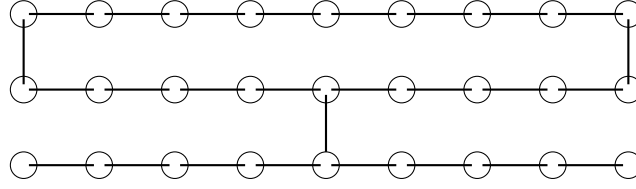


Figure 1.3: An example of two cluster states connected by a single vertical link at the center node, representing a two-qubit entangling operation between the clusters.

1.5 Measurement-based quantum computation

Before delving into the main subject of this thesis, it is appropriate to briefly discuss an alternative paradigm for quantum computation: the cluster-state model. In fact, this seems to be the most promising route in PQC, nowadays, as compared to the more “old fashion” LOQC approach. In fact, cluster states offer an alternative paradigm to the traditional circuit-based model of quantum computing, which is known as *one-way quantum computing* or *measurement-based quantum computing* (MBQC), originally introduced by Raussendorf and Briegel [57].

In essence, the MBQC paradigm proposes that the quantum computation is driven entirely by sequential single-qubit measurements performed on a pre-entangled resource state, rather than applying sequences of single- and two-qubit gates to evolve an initially given quantum state. A *cluster state* is a highly entangled multi-qubit state, in which physical qubits are represented as nodes in a lattice, and entanglement between neighbouring qubits is represented by edges, as illustrated in Fig. 1.3. Formally, a cluster state $|C\rangle$ is defined as an eigenstate of a set of commuting operators S_i , called the stabilizer generators, one for each qubit i in the cluster [58]:

$$S_i |C\rangle = \pm |C\rangle \forall i$$

and

$$S_i = X_i \bigotimes_{j \in \mathcal{N}(i)} Z_j$$

where X_i and Z_j are Pauli operators and $\mathcal{N}(i)$ denotes the set of nearest neighbours of qubit i .

Equivalently, a cluster state can be prepared by initializing all qubits in the $|+\rangle \equiv (|0\rangle + |1\rangle)/\sqrt{2}$ state and then applying a controlled-phase (CPHASE) gate between every pair of neighbouring qubits [59].

With the cluster state prepared, computation proceeds by measuring qubits one column at a time from left to right. Single-qubit gates are implemented by choosing an appropriate measurement basis for each qubit, while two-qubit entangling gates arise naturally from the vertical links connecting qubits in adjacent rows [59]. The outcome of each measurement is classically feedforward to determine the basis of the next measurement, so that the overall computation is steered adaptively. A computational basis measurement simply removes a qubit from the cluster, leaving behind a known local Z^m correction on neighbouring qubits that can be compensated by classical post-processing [59]. Crucially, two-dimensional cluster states are necessary and sufficient for universal quantum computation [58]: it has been shown that any quantum circuit of breadth n and depth $d(n)$ can be efficiently simulated on a two-dimensional cluster state [57],

whereas linear cluster chains are classically simulable and therefore insufficient [59]. In the context of LOQC, this approach is particularly attractive since single-qubit measurements on photons are experimentally straightforward and naturally replace the need for challenging deterministic two-qubit gates. However, generating the required entangled cluster state with linear optics and probabilistic gates remains the central difficulty: building the cluster requires a significant overhead in photon number and entangling operations, and photon loss remains a critical challenge that must be addressed through suitable error-correction strategies [25]. The Ascella platform previously introduced is currently being explored also in the direction of MBQC, although only proof-of-principle gate demonstrations on few qubits have been reported, so far [24].

Summary

In this broad introductory Chapter we have laid the theoretical and physical foundations necessary to grasp the developments of this thesis, which will follow in the next Chapters. Starting from the basic formalism of the circuit-model of quantum computation, we have introduced the main physical platforms that are currently explored to realize this computational paradigm, with particular attention to the photonic approach. We have seen that photons are natural carriers of quantum information, but the absence of photon–photon interactions poses a fundamental obstacle to the implementation of deterministic two-qubit gates.

Within this context, the KLM scheme offers an elegant probabilistic resolution to the latter issue, exploiting quantum interference and projective measurements as an effective non-linearity. On an alternative direction, we have also briefly mentioned measurement-based paradigms, such as cluster-state computation, which represent a promising route to realize quantum photonics computing in practice, although it is one that falls outside the scope of this thesis.

In the following Chapter 2, we will start from a rigorous treatment of the KLM protocol, then we will develop the full theoretical framework underlying important gates made by linear optical quantum computation, and we will present in detail the explicit construction of a CNOT gate realized with six waveguide channels, the central result around which this thesis is built.

Chapter 2

Two-qubit gates in linear optical quantum computing

The general scheme of Linear Optical Quantum Computing (LOQC) offers several advantages over competing approaches to quantum computing, as already introduced in the previous Chapter. First, the basic optical elements do not require cryogenic temperatures, unlike many alternative solid-state proposals. Second, photons naturally maintain their coherence over timescales that are long compared to the elementary control operations, reducing the impact of decoherence. Third, dominant noise sources are better understood and do not depend on thermal interactions that are difficult to predict or characterize. The resource overhead per reliable logical gate is large but not excessive, and is comparable to that of other leading proposals for scalable quantum computation. This puts LOQC on an equal footing with alternative approaches, making it a genuinely competitive technology for scalable quantum information processing.

In the previous chapter, we have introduced the general principles of Photonic Quantum Computing and its main features. In particular, we have already discussed how large cross-Kerr nonlinearities can, in principle, mediate a single-photon CNOT gate. We have also seen that, unfortunately, nonlinearities of this magnitude are far beyond what nature provides — they fall short by many orders of magnitude. Even if such a material is used to create long fibers for enhanced nonlinearities, photon losses in the Kerr medium will prevent the gate from operating properly. Furthermore, Kerr nonlinearities are typically very noisy. LOQC thus represents a most promising alternative, relying on projective measurements performed by photodetectors to engineer an effective photon–photon interaction. However, the drawback of this approach is that the resulting optical quantum gates are inherently probabilistic: the gate fails with non-negligible probability, thus corrupting the encoded quantum information. Although this issue can be worked around by scaling up the number of optical modes exponentially, practical scalability demands that the required modes grow only polynomially.

A breakthrough in LOQC occurred when Knill, Laflamme, and Milburn (KLM) introduced a protocol that bypasses this limitation, proving that linear optics, ancilla photons, and feed-forward measurements suffice for scalable universal quantum computation [5], albeit probabilistic (as intrinsically constrained by the presence of only linear optical elements). In the following years, their result triggered an extensive research effort along two complementary directions. On one hand, a substantial theoretical literature investigated the fundamental limits of post-selected linear-optical gates,

sharpening and eventually closing the bounds on their achievable success probabilities [60, 6, 7, 8, 61]. On the other hand, several groups pursued dedicated, resource-frugal constructions of the CNOT gate itself, trading the asymptotic scalability of the KLM architecture for immediate experimental feasibility with few-photon sources [9, 10, 11, 12]. Among these, the scheme proposed by Ralph [10] is of particular interest, as it dispenses entirely with ancilla photons and was the first to be experimentally realized, both in bulk optics [13] and, later, in an integrated photonic platform [14]. We will review these protocols in detail, in the following of this Chapter.

2.1 The KLM protocol

In 2001, Knill, Laflamme, and Milburn [5] proved that universal quantum computation is possible using only linear optical elements such as beam splitters and phase shifters, in addition to single-photon sources and single-photon detectors. Prior to this result, the main obstacle to scalable optical quantum information processing (QIP) had been the apparent need for nonlinear couplings between optical modes, a requirement that is technically demanding to satisfy at the few-photon level.

In the KLM scheme, qubits are each encoded in a single photon occupying one of two optical modes. The scheme exploits feedback from photo-detectors and quantum teleportation to inject probabilistic two-photon gates into a quantum circuit with arbitrarily high success probability, and leverages quantum error correction to suppress the residual failure rate to fault-tolerant levels [62][63].

The sufficiency of linear optics for efficient QIP is established through three key results. First, non-deterministic quantum computation is shown to be possible: a probabilistic nonlinear sign shift, acting on one qubit and using two ancilla photons with post-selection, succeeds with probability $1/4$, and from two such operations a conditional sign flip between two qubits can be constructed, succeeding with overall probability $1/16$. These elementary gates rely on post-selection alone, with no feedback involved: the ancilla measurements outcome simply determines whether the operation is accepted or discarded. Second, the success probability of the resulting two-qubit gate can be boosted arbitrarily close to unity by consuming entangled ancilla states prepared non-deterministically and applying quantum teleportation with feedback: here, the measurement outcome determines which correction to apply, rather than whether to discard the result. Third, by employing quantum codes (in particular, a two-qubit stabilizer code χ_2 that can be concatenated), the resource overhead for achieving accurate logical qubits scales efficiently with the target accuracy, completing the argument for efficient LOQC.

Taken together, these results imply that LOQC can be implemented robustly with polynomial resources, placing it on a comparable footing with other leading proposals for scalable quantum computation.

Basic elements of the scheme. The elementary building blocks of the KLM scheme are bosonic qubits, each encoded in a single photon distributed across two optical modes, with the logical states $|0\rangle$ and $|1\rangle$ corresponding to the photon occupying one mode or the other. Single-qubit rotations are implemented straightforwardly via phase shifters and beam splitters, which together can realize any unitary transformation on the mode

pair. Readout is performed by photo-detectors, which destructively measure the photon number in a given mode; an approximate photon counter, sufficient for the purposes of LOQC, is constructed by splitting the mode across N output ports and placing one detector at each. The remaining challenge, which constitutes the core of the KLM proposal, is the realization of a non-trivial two-qubit gate such as the conditional sign flip, also known as CZ gate that we have already seen in Chapter 1, in particular in Sec. 1.3.3.

Quantum teleportation. The basic teleportation protocol (T_1) transfers the state $\alpha_0|0\rangle_1 + \alpha_1|1\rangle_1$ of mode 1 to mode 3. The protocol begins by preparing the ancilla entangled state $|t_1\rangle_{23} \propto |01\rangle_{23} + |10\rangle_{23}$, and adjoining it to the input. Modes 1 and 2 are then jointly measured in the Bell basis $\{|01\rangle_{12} \pm |10\rangle_{12}, |00\rangle_{12} \pm |11\rangle_{12}\}$, which proceeds in two steps: first, a parity measurement determines whether the total photon number in modes 1 and 2 is odd or even, and then the sign of the superposition is resolved.

If the parity is odd and the sign is $+$, mode 3 is left in the state $\alpha_0|0\rangle_3 + \alpha_1|1\rangle_3$, recovering the input exactly; if the sign is $-$, a phase shifter corrects the state to $\alpha_0|0\rangle_3 + \alpha_1|1\rangle_3$. If the parity is even, the roles of $|0\rangle$ and $|1\rangle$ are exchanged in mode 3, and this cannot be corrected by linear optical means alone.

In linear optics, the Bell measurement is implemented by applying a balanced beam splitter to modes 1 and 2, and then counting the output photons, which determines the parity; when the parity is odd, also the sign. This partial Bell measurement (BM_1) succeeds with probability $1/2$, and upon success the input state is faithfully transferred to mode 3 up to a known and correctable phase.

The success probability of the teleportation protocol T_1 can be arbitrarily boosted close to unity by generalising the ancilla entangled state. Specifically, one defines a family of states $|t_n\rangle$, living in the space of n bosonic qubits, such that applying the $(n+1)$ -point Fourier transform $\hat{F}_{n+1}$ followed by photon-number measurement on $n+1$ modes realises a generalised Bell measurement BM_n . The outcome of this measurement is the total number of detected photons $k \in \{0, 1, \dots, n+1\}$. If $0 < k < n+1$, teleportation succeeds: the input state reappears in the output mode $n+k$ up to a correctable phase shift, and the remaining modes are returned in the state $|1\rangle$, so that their photons can be reused in future operations. The only two failure cases are $k=0$ and $k=n+1$, in which the measurement collapses the input into $|0\rangle$ or $|1\rangle$ respectively, destroying the superposition. Each failure event occurs with probability $1/(n+1)$ regardless of the input state, so that a single teleportation succeeds with probability $n/(n+1)$.

This near-deterministic teleportation directly enables a conditional sign flip $c\text{-}z_{n^2/(n+1)^2}$. The key idea is to reduce the implementation of the two-qubit gate to a state preparation problem: rather than applying $c\text{-}z$ directly to the two qubits, one prepares offline a specially entangled resource state $|\text{cs}_n\rangle$ that encodes the action of the gate, and then teleports the first mode of each qubit through it. The measurement outcomes of the two teleportations extract the effect of the gate on the output, without the two qubits ever interacting directly. Since the gate requires two independent teleportations, the overall success probability is:

$$p_{\text{succ}} = \left(\frac{n}{n+1} \right)^2 = \frac{n^2}{(n+1)^2}.$$

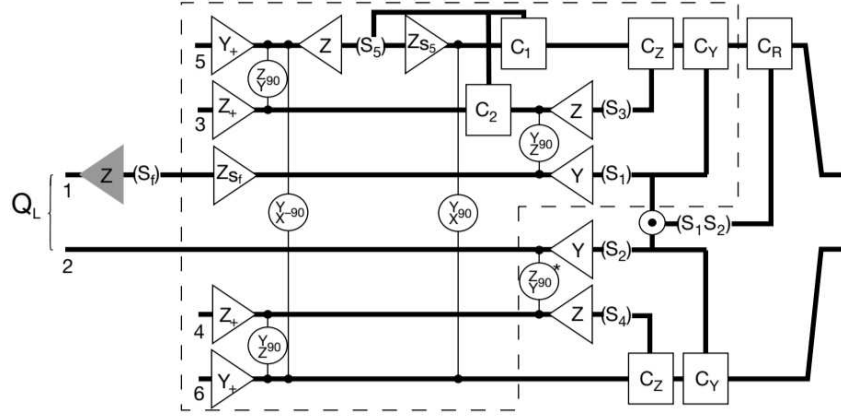


Figure 2.1: Recovery from Z measurement. From [5].

The preparation of $|cs_n\rangle$ can be attempted repeatedly and in parallel, offline, without stalling the main circuit, so that the non-determinism is effectively confined to a classical resource overhead.

The two-qubits quantum codes χ_2 . At this stage of the KLM protocol, the gate $c\text{-}Z_{n^2/(n+1)^2}$ can be made to succeed with probability arbitrarily close to unity, but at the cost of preparing increasingly complex resource states $|cs_n\rangle$, which becomes prohibitively demanding for large n . The solution proposed by the authors is to exploit quantum error-correcting codes to achieve the same goal far more efficiently, keeping n small and instead iterating a coding procedure.

The central tool is a *two-qubit stabilizer code* χ_2 , which encodes a single logical qubit into two physical bosonic qubits via the logical basis states:

$$|0\rangle_L = |00\rangle + |11\rangle \quad |1\rangle_L = |01\rangle + |10\rangle$$

The logical Pauli operators X_L , Y_L and Z_L are implemented by combining physical operations on the two constituent qubits. Logical 180° rotations can always be executed without failure by directly applying the corresponding physical rotations. The critical operations are the logical 90° rotations, such as $Z_L^{90^\circ}$ and $(Z_{L_1}Z_{L_2})^{90^\circ}$, which are implemented via the teleportation networks and can fail.

However, the code is specifically designed so that a failure results at worst in an unintended Z measurement of the logical qubit, a structured and detectable error that can be corrected by the recovery network (Figure 2.1). With this failure protocol, the logical failure probability f_z satisfies

$$f_z < \mathcal{O}(f^2)$$

and in particular $f_z < f$ whenever $f < 1/6.43$, where $f = 1/(n+1)$ is the physical gate failure probability. This quadratic suppression of the error is achieved without requiring large n , since even modest values are sufficient to place f below the threshold $1/6.43$. The threshold $f < 1/6.43$ is not chosen arbitrarily, but emerges naturally from the quantitative analysis of the recovery protocol shown in Fig. 2.1, which bounds the residual failure probability f_r as a function of the physical failure rate f .

The efficiency of the scheme is further improved by concatenation: by treating logical qubits as the physical qubits of a higher-level encoding and iterating the procedure, the logical failure probability is suppressed doubly exponentially in the number of concatenation levels. This is the standard mechanism underlying fault-tolerant quantum computation, and it implies that arbitrarily accurate logical gates can be achieved with only a polynomial overhead in physical resources. Concretely, the authors show that just two levels of χ_2 concatenation are sufficient to bring the logical failure probability below 5%, requiring approximately 300 successful $c\text{-}z_{9/16}$ operations per logical two-qubit gate. This establishes that LOQC is not only possible in principle, but genuinely scalable in practice.

2.2 Theoretical bounds for post-selected gates

As discussed above, the KLM scheme relies on two conceptually distinct ingredients: elementary post-selected gates (NS and CZ), which involve no feedback, and a teleportation-based architecture that uses feedback to boost the overall two-qubit gate to near-unit success probability. The latter is quite challenging to implement, with state-of-the-art technology. Hence, the bounds discussed in this section concern only the first ingredient: the maximum success probability achievable by NS and CZ themselves, within the postselection-only (no-feedback) framework. A natural question arises: given that these gates are inherently probabilistic, what is their maximum achievable success probability? We believe this is the most meaningful question to answer, in view of practically realizing even proof-of-principle demonstrations of the KLM protocol.

2.2.1 Improving the CZ gate

In 2002, Knill improved on the original KLM construction, reducing both the circuit complexity and increasing the success probability of the CZ gate [60]. In the KLM scheme, the CZ gate is obtained by composing two independent NS gates, requiring six optical modes, six beam splitters, two photon counters and two additional photodetectors, with overall success probability set to $1/16$, as already introduced previously. In the revised approach by Knill, the CZ gate is implemented directly as a single optimized linear optics network on four modes, using four beam splitters, two initial phase shifts, two helper photons and two photon counters, which allows to reach a success probability of $2/27$.

The value $2/27$ is obtained through a combination of algebraic and numerical methods. The action of the gate is governed by the 4×4 submatrix V of the unitary U associated with the linear optics transformation. The entries of V must satisfy polynomial identities ensuring correct gate behaviour on all input states; these are first solved analytically, exploiting scaling symmetries to reduce the number of free parameters. The remaining variables are then optimized numerically in MATLAB, maximizing the probability of success — defined as the squared amplitude for the ancilla modes to return to their initial state — subject to the constraint that V extends to a unitary matrix. This consistently converges to the matrix V_{180° , with success probability $2/27$.

Beyond this improved construction, Knill also shows that a success probability of 1 is fundamentally unattainable for the CZ gate using only passive linear optics and post-

selection. The argument relies on the fact that the states producible in the output modes, after tracing out the helper modes, are necessarily mixtures of products of linear expressions in the creation operators. The target Bell state cannot be written in this factored form, since the associated polynomial is irreducible, ruling out perfect success regardless of the number of helper photons.

Interestingly, this leaves open the question of how far from 1 the true optimum lies, which Knill explicitly poses as an open problem in LOQC: what is the maximum success probability for CS_θ using passive linear optics with k helper photons and postselection without feedback? For $\theta = 180^\circ$ and $k \geq 2$, $2/27$ is achieved, but its optimality remained actually unknown.

2.2.2 Quantitative bounds

This gap was later addressed by Knill's 2003 Brief Report, in which he develops a general technique for upper-bounding the success probability of postselected gates implemented with single-photon sources, passive linear optics, and photon counting without feedback. The key tool behind these bounds is simple: if you start from the vacuum, add some single helper photons, and apply only passive linear optics (beam splitters and phase shifters), the resulting state can never have, on average, more than one photon in any given mode. This is essentially because passive linear optics only redistributes photons among modes, it cannot concentrate them.

Knill exploits this fact by using the NS and CZ gates as building blocks inside a larger optical network designed to produce a state with an unusually high average photon number in one mode, higher than what linear optics alone could ever achieve.

NS bound. For the NS gate, the idea works as follows. Suppose NS succeeds with probability p . One can build a network that, conditioned on this success, produces the two-photon state $|2\rangle_a$ in some mode a . Before the final postselection, the state of mode a is therefore a mixture

$$\rho = p |2\rangle\langle 2| + (1 - p) \rho',$$

whose average photon number is $2p + x$, with $x \geq 0$ coming from ρ' . But ρ itself is just a state produced by linear optics acting on single photons, so it must obey the bound above: its average photon number cannot exceed 1. This immediately gives

$$2p \leq 1 \quad \implies \quad p \leq \frac{1}{2},$$

which is the bound $p_{\max}(\text{NS}) \leq 1/2$.

CS bound. The argument for CZ follows exactly the same logic, but the construction is slightly more involved. Using a sequence of beam splitters together with one application of CZ (success probability p), one prepares a state in which a certain logical mode: a specific linear combination of three physical modes has an average photon number of $4/3$. Again, before post-selection, this contribution only appears with probability p , so the overall average photon number in that mode is at least $p \cdot 4/3$. Since this must still be ≤ 1 , one obtains

$$p \cdot \frac{4}{3} \leq 1 \quad \implies \quad p \leq \frac{3}{4},$$

giving $p_{\max}(\text{CS}) \leq 3/4$.

In both cases, the strategy is the same: use the gate to artificially overload a mode with more photons than linear optics could ever produce on its own. Since this overloaded configuration occurs only in the successful branch of the gate, the universal bound on the average photon number directly translates into an upper bound on the success probability p .

These bounds apply to any implementation within the LOP+PC framework (a procedure using single helper photons and linear optics + post-selection based on measured photon counts), regardless of the specific optical network used. This is a fundamentally different kind of result from the constructive results of [5] and [60]. However, the gap between these bounds and the best known constructions remains large, particularly for CZ, where $2/27$ sits far below the bound of $3/4$, leaving the determination of the true optimum an open problem.

While Knill's bound $p_{\max}(\text{NS}) \leq 1/2$ left a significant gap with respect to the achievable value of $1/4$, this gap was closed in subsequent work. Scheel and Lütkenhaus [7] introduced an abstract model for linear-optics networks, reducing any such network to a single active beam splitter connecting the signal mode to the ancilla. Analyzing broad classes of ancillary states, they found semi-analytical and numerical evidence that the success probability of NS never exceeds $1/4$, conjecturing this to be the general bound. This conjecture was proven rigorously by Eisert [8], who linked the optimization of success probabilities to convex optimization and Lagrange duality, obtaining bounds valid for postselected networks of arbitrary size and ancilla dimension. That solved the difficulty, left open since Knill's work, of bounding the network size a priori. The result establishes that $p_{\max}(\text{NS}) = 1/4$ exactly, without feedforward: not merely an upper bound, but the true optimum, already saturated by the original KLM construction. The analogous question for CZ remains open: no comparably tight bound has been established, and the gap between the known construction ($2/27$, [60]) and Knill's bound ($3/4$, [6]) is still unresolved.

Having established $1/4$ as the optimal success probability for NS without a feedforward scheme, a further natural question then arises: how much feedforward information, i.e., how much the use of the outcome of a failed attempt to apply a corrective network, can improve this value? Knill's bound of $1/2$ [6] allows in principle for a significant improvement, but does not address whether this bound is saturated. This question was investigated by Scheel [61], who analysed the three-mode KLM implementation of NS and classified its failure outcomes into irrecoverable errors, where information about the signal state is destroyed, and recoverable errors, which can in principle be corrected by a further conditional network. Applying a single round of such corrective feedforward raises the success probability only marginally, from $1/4$ to ≈ 0.272 , with multiple rounds reaching at most ≈ 0.28 . Moreover, the maximal failure probability obtained (≈ 0.66) suggests that the true bound with feedforward is much tighter than Knill's $1/2$, closer to $1/3$. These results indicate that, at the level of the elementary NS gate, the benefit of feedforward is too small to justify its cost in additional resources.

2.3 Implementations of linear-optics CNOT

The constructions discussed so far share a common architectural philosophy: the elementary NS gate is treated as a modular primitive, whose success probability ($1/4$ at best, without feed-forward [8]) is then amplified, via teleportation and quantum coding, to near-unit probability at the level of the logical CZ/CNOT gate. In this view, the probability of success for the NS gate is not the end goal but the input to a larger machine designed for asymptotic scalability.

The works that we will discuss in this section take a clearly different approach. Rather than composing CZ/CNOT out of NS gates, they design a single, dedicated optical network that implements the two-qubit gate directly, exploiting the specific structure of polarization-encoded qubits and entangled (or single-photon) ancillae. There is no reusable "NS module" here: each network is tailored ad hoc to the CNOT operation itself, and the resulting success probabilities are constant numbers, not starting points for further amplification via teleportation. The goal is not asymptotic scalability, but immediate experimental feasibility with few-photon sources and detectors.

A second, related difference concerns the role of feed-forward. In the NS context, feed-forward meant attempting to recover from a failed gate by applying an additional conditional optical network, with strictly marginal gains ($0.25 \rightarrow 0.28$, [61]). Here, feed-forward instead means accepting additional successful detection outcomes that differ from the canonical one only by a local, deterministic single-qubit correction (a phase flip or a bit flip on the output qubits). Since such corrections cost nothing in linear optics, this effectively doubles the success probability "for free" by exploiting a symmetry of the network, rather than trying to rescue an otherwise lost outcome.

Taken together, these proposals (see, in particular, Refs. [9, 10, 11, 12]) trace the progression from theoretical proposal to the first experimental demonstrations of a photonic CNOT gate, illustrating a complementary, resource-frugal path to linear-optical quantum logic alongside the KLM/Knill program.

2.3.1 CNOT via quantum erasure

The theoretical proposal. Pittman, Jacobs and Franson [9] proposed a CNOT gate built entirely from polarizing beam splitters (PBS) that fully transmit one polarization and fully reflect the orthogonal one, without relying on quantum teleportation. Qubits are encoded in the polarization of single photons, with $|H\rangle \equiv |0\rangle$ and $|V\rangle \equiv |1\rangle$, and measurements are performed both in this basis and in the rotated basis:

$$|F\rangle = \frac{1}{\sqrt{2}}(|H\rangle + |V\rangle)$$

$$|S\rangle = \frac{1}{\sqrt{2}}(|H\rangle - |V\rangle)$$

The construction, schematically described in Fig. 2.2 proceeds by first defining three elementary operations, each implemented with a single PBS followed by a single polarization-sensitive detector, and only afterwards combining them into a full CNOT.

The *quantum parity check* takes an arbitrary input qubit

$$|\text{in}\rangle_{2'} = \alpha |H\rangle_{2'} + \beta |V\rangle_{2'}$$

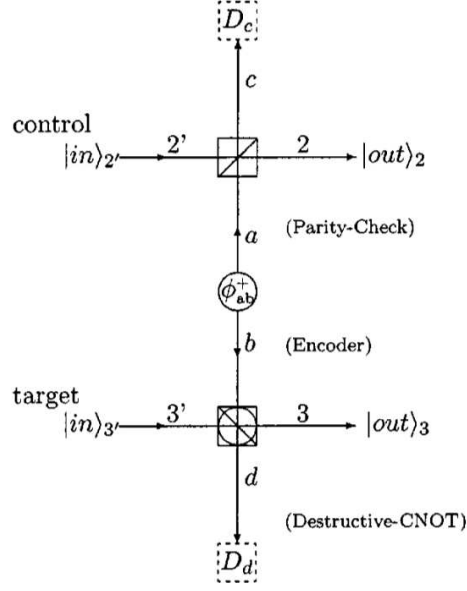


Figure 2.2: Construction of a probabilistic CNOT gate by combining a quantum encoder with a destructive CNOT. Original picture taken from Ref. [9].

together with a second photon in the diagonal state

$$|\phi\rangle_a = \frac{1}{\sqrt{2}}(|H\rangle_a + |V\rangle_a)$$

and mixes them on a PBS. Writing the resulting amplitudes in the FS basis for the detection mode c , one finds that conditioning on detecting a single F -polarized photon (and none in S) maps the input state, unchanged, onto the output mode 2, with probability $1/4$. Crucially, the complementary detection outcome (a single S -polarized photon) maps the input onto the phase-flipped state $\alpha|H\rangle - \beta|V\rangle$. Since this differs from the desired output only by a known, local single-qubit correction, accepting this outcome as well (and applying the corresponding π phase shift) raises the success probability from $1/4$ to $1/2$ at no additional resource cost.

The *destructive CNOT* follows the same pattern, but with the two input photons mixed on a PBS oriented in the FS basis: a target photon

$$|in\rangle_{3'} = \alpha|H\rangle_{3'} + \beta|V\rangle_{3'}$$

and a control photon in a definite polarization state. If the control is $|V\rangle$, the surviving detection outcomes correspond to a flip of the target ($H \leftrightarrow V$); if the control is $|H\rangle$, the target is left unchanged. As before, one of the two acceptable detector outcomes requires no correction, while the other requires a local single-qubit operation (a 90° polarization rotation followed by a conditional π phase shift); accepting both raises the success probability from $1/4$ to $1/2$. The price paid for this operation is that the control photon is consumed in the process, hence the name destructive.

Finally, the *quantum encoder* is structurally identical to the parity check, but one of its two inputs is now one half of a maximally entangled pair

$$|\phi_{ab}^+\rangle = \frac{1}{\sqrt{2}}(|H_a H_b\rangle + |V_a V_b\rangle)$$

A successful detection event — with probability $1/2$, including the feed-forward correction as above — post-selects the transformation

$$\alpha|H\rangle_{2'} + \beta|V\rangle_{2'} \rightarrow \alpha|H_2H_b\rangle + \beta|V_2V_b\rangle$$

that is, the value of the input qubit is copied onto the second output mode b while being preserved in mode 2. This is not in tension with the no-cloning theorem, since the operation does not clone an arbitrary unknown state for free: it consumes one half of a pre-shared entangled pair as a resource.

Combining these elements yields a full, non-destructive CNOT (see Fig. 2.2): the quantum encoder first duplicates the control qubit, sending one copy directly to the output and the other into the destructive CNOT together with the target qubit. The overall success probability is the product of the two stages:

$$p_{\text{succ}} = \frac{1}{2} \cdot \frac{1}{2} = \frac{1}{4},$$

requiring only a single two-photon entangled ancilla pair $|\phi_{ab}^+\rangle$.

A key physical ingredient throughout this scheme is “quantum erasure”: naively, the photons reaching the detectors carry which-path information that would reveal the polarization of the input qubits, destroying the coherence of any superposition state. By measuring these detectors in the rotated FS basis rather than HV , this information is erased, because an F -polarized photon is an equal superposition of H and V , and its detection therefore reveals nothing about the original qubit values. This is the mechanism that allows the gate to act coherently on arbitrary superposition inputs, rather than merely on computational basis states.

The authors explicitly contrast this result with NS-based constructions of CZ, which require the same number of ancilla photons but achieve only $p = 1/16$ [5]. This highlighting that a direct, dedicated network can outperform the NS-based modular approach at equal resource cost, at the price of abandoning the asymptotic scalability that the NS/teleportation architecture provides.

Experimental realization. The CNOT scheme proposed in Ref. [9] was experimentally realized by Gasparoni *et al*, who reported the first demonstration of a CNOT gate for independent photons satisfying the feed-forwardability criterion [12]. The experiment implements only the passive version of the gate, without the local corrections described above: both the quantum encoder and the destructive CNOT are operated accepting only their primary detection outcome, each succeeding with probability $1/4$ rather than $1/2$. The overall theoretical success probability is therefore reduced from $1/4$ to $1/16$, further limited experimentally by a fidelity factor of ≈ 0.79 . Despite this reduced rate, the authors verified correct CNOT operation on all logical basis inputs and demonstrated the generation of entanglement from separable inputs, confirming that the gate performs a genuinely coherent two-qubit operation rather than a classical logic function.

2.3.2 CNOT with a single ancilla photon

Pittman, Fitch, Jacobs and Franson [11] addressed a practical limitation of the scheme in [9]: the reliable generation of entangled ancilla pairs remained experimentally challenging. They proposed and demonstrated a simplified CNOT gate (Fig. 2.3) in which

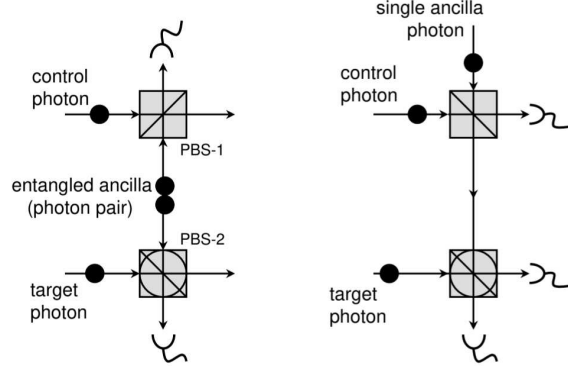


Figure 2.3: Comparison of the two implementations of a CNOT gate using linear optics and ancilla photons proposed by Pittman. On the left the gate proposed in Ref. [9], with two entangled ancilla photons. On the right the gate proposed in Ref. [11], which requires only one ancilla photon. PBS-1 and PBS-2 are polarizing beam splitters, with PBS-2 being rotated by 45° . Original picture taken from Ref. [11].

the entangled ancilla pair of the quantum encoder is replaced by a single ancilla photon prepared in the equal superposition $\frac{1}{\sqrt{2}}(|0\rangle_A + |1\rangle_A)$. This single photon, fed into the upper beam splitter where a detector previously sat, still allows the value of the control qubit to be copied onto the two output ports, reproducing the role of the entangled-pair encoder with half the ancilla resources.

This simplification comes at a cost. With an entangled ancilla pair, the success of the gate could be certified by measuring the ancilla photons alone, independently of the logical qubits (feed-forwardability). With a single ancilla photon, the isomorphism between the input and output Hilbert spaces requires that exactly one photon emerge from each of the three output ports — ancilla, control, and target — a condition that, without quantum non-demolition detectors, can only be verified by eventually detecting the logical output qubits themselves. The gate therefore operates in the coincidence basis, sacrificing feed-forwardability for a reduction in ancilla resources.

Explicitly, for an ancilla in $\frac{1}{\sqrt{2}}(|0\rangle_A + |1\rangle_A)$ and an arbitrary two-photon control-target state $\alpha_1 |0_c 0_t\rangle + \alpha_2 |0_c 1_t\rangle + \alpha_3 |1_c 0_t\rangle + \alpha_4 |1_c 1_t\rangle$, the output state takes the form

$$\begin{aligned}
 |\psi\rangle_{\text{out}} = & \frac{1}{2\sqrt{2}} |0_A\rangle (\alpha_1 |0_c 0_t\rangle + \alpha_2 |0_c 1_t\rangle + \alpha_3 |1_c 1_t\rangle + \alpha_4 |1_c 0_t\rangle) + \\
 & + \frac{1}{2\sqrt{2}} |1_A\rangle (\alpha_1 |0_c 1_t\rangle + \alpha_2 |0_c 0_t\rangle + \alpha_3 |1_c 0_t\rangle + \alpha_4 |1_c 1_t\rangle) + \\
 & + \frac{\sqrt{3}}{2} |\psi_\perp\rangle,
 \end{aligned}$$

where $|\psi_\perp\rangle$ collects all amplitudes incompatible with the coincidence condition of one photon in each of the three output modes.

The first term is precisely the CNOT transformation of the input, obtained with probability $1/8$ upon detecting the ancilla in $|0\rangle_A$. The second term corresponds to the CNOT composed with an additional, unwanted bit flip on the target ($0_t \leftrightarrow 1_t$); conditioned on detecting the ancilla in $|1\rangle_A$, this outcome can in principle be corrected by a feed-forward bit flip, raising the overall success probability from $1/8$ to $1/4$. In the experimental demonstration briefly described in the following, this correction was not

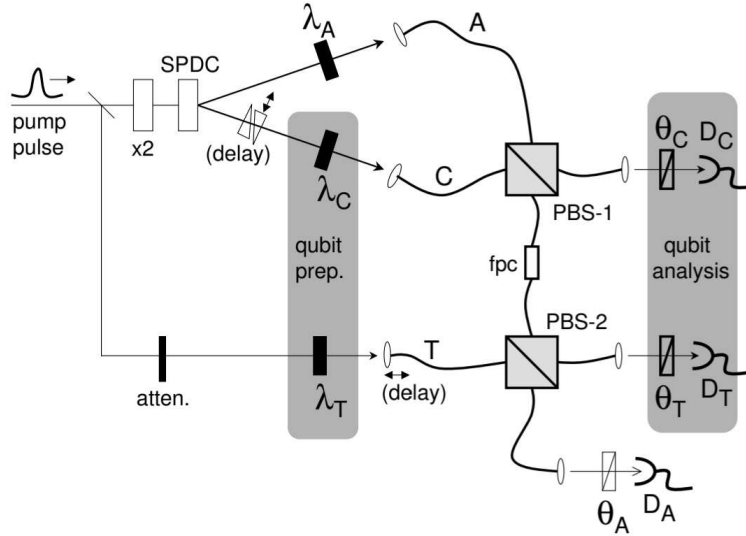


Figure 2.4: Experimental apparatus used to demonstrate the CNOT gate. Original picture taken from Ref. [11].

implemented, and only the $|0\rangle_A$ outcome was used.

Experimental realization. The experimental apparatus employed for the CNOT demonstration is shown in Fig. 2.4. In the experiment, spontaneous parametric down-conversion was used for producing the control and ancilla photons; the target photon was instead obtained from a strongly attenuated portion of the original laser pulse. All the three photons were combined at two polarizing beam splitters, with single-mode fibres, narrow-band filters, and path-length tuning used to maximise their mutual indistinguishability. Measuring all possible combinations of logical inputs $\{0, 1\}^2$, the device reproduced the truth table of a CNOT gate up to technical errors of order 20%. Feeding in a control qubit in the superposition $\frac{1}{\sqrt{2}}(|0\rangle + |1\rangle)$ with the target in $|0\rangle$, the output was found to be entangled, with a visibility of $(61.5 \pm 7.4)\%$ in the $|\phi^+\rangle$ Bell state. This represented the first experimental demonstration of a CNOT gate for independent photons, albeit in a non-scalable, single-ancilla, coincidence-basis configuration.

2.4 Probabilistic CNOT without ancilla photons

All the schemes discussed so far require ancilla photons — i.e., either an entangled pair [9, 12] or a single prepared photon [11] — as a resource to mediate the effective photon–photon interaction. A natural question is whether ancilla photons are strictly necessary, or whether the two logical input photons alone, combined through a purely linear-optics network, can already produce the desired CNOT transformation via postselection.

In fact, Ralph et al. proposed a CNOT gate constructed from a minimal circuit architecture made of only five beam splitters and no ancilla photons [10], operating in the coincidence basis with a maximum success probability of $1/9$. This has represented a ground-breaking result, as we will discuss, and it has been implemented in different experimental configurations and set-ups.

Essentially, the proposed gate consists of five beam splitters arranged in the network

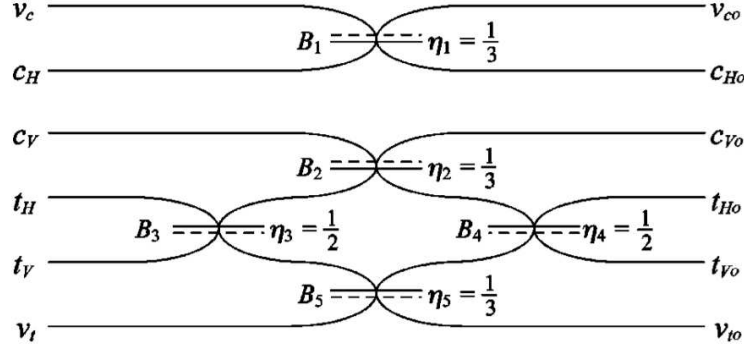


Figure 2.5: Schematic representation of the circuit implementing the coincidence CNOT gate according to Ralph et al. [10]. Dashed lines indicate the surface from which a sign change occurs upon reflection. The control modes are c_H and c_V . The target modes are t_H and t_V , respectively. Modes v_c and v_t are unoccupied ancillary modes. This picture is taken from the original Ref. [10].

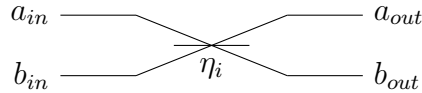
as schematically shown in Fig. 2.5: beam splitters B_1 , B_2 , and B_5 have reflectivity $\eta_1 = \eta_2 = \eta_5 = 1/3$, respectively, while B_3 and B_4 have reflectivity $\eta_3 = \eta_4 = 1/2$. All beam splitters are assumed asymmetric in phase. Reflection off the bottom produces the sign change, except for B_1 and B_2 which have a sign change by reflection off the top. So, the input-output relations (Heisenberg equations), between two input mode operators (a_{in} and b_{in}) and the corresponding output operators (a_{out} and b_{out}), for the B_3, B_4, B_5 beam splitters have the form:

$$\begin{aligned} a_{out} &= \sqrt{\eta_i} a_{in} + \sqrt{1 - \eta_i} b_{in} \\ b_{out} &= \sqrt{1 - \eta_i} a_{in} - \sqrt{\eta_i} b_{in} \end{aligned} \quad (2.1)$$

For B_1, B_2 beam splitters have the form:

$$\begin{aligned} a_{out} &= -\sqrt{\eta_i} a_{in} + \sqrt{1 - \eta_i} b_{in} \\ b_{out} &= \sqrt{1 - \eta_i} a_{in} + \sqrt{\eta_i} b_{in} \end{aligned} \quad (2.2)$$

where η_i is the reflectivity of the beam splitter, as already introduced in Chapter 1.



The authors assumed polarizing qubits in dual-rail encoding. The control qubit is encoded in the modes c_H, c_V and the target in t_H, t_V ; v_c and v_t are two additional vacuum modes, unoccupied input ports of the beam splitters. We notice that the latter should not be confused with ancilla photons: they require no additional photon source or state preparation, and carry no resources beyond the two logical input photons themselves.

Analytic derivation of the optimal beam splitting ratios. Let us now develop the evolution of a general input within the scheme, using the beam splitting operations parametrized in Eqs. 2.1 and 2.2:

$$\left\{ \begin{array}{l} c_{Ho} = \sqrt{1 - \eta_1} v_c + \sqrt{\eta_1} c_H, \\ c_{Vo} = -\sqrt{\eta_2} c_V + \sqrt{1 - \eta_2} (\sqrt{\eta_3} t_H + \sqrt{1 - \eta_3} t_V), \\ t_{Ho} = \sqrt{\eta_4} (\sqrt{1 - \eta_2} c_V + \sqrt{\eta_2} \sqrt{\eta_3} t_H + \sqrt{\eta_2} \sqrt{1 - \eta_3} t_V) + \\ \quad + \sqrt{1 - \eta_4} (\sqrt{\eta_5} \sqrt{1 - \eta_3} t_H - \sqrt{\eta_5} \sqrt{\eta_3} t_V + \sqrt{1 - \eta_5} v_t), \\ t_{Vo} = \sqrt{1 - \eta_4} (\sqrt{1 - \eta_2} c_V + \sqrt{\eta_2} \sqrt{\eta_3} t_H + \sqrt{\eta_2} \sqrt{1 - \eta_3} t_V) - \\ \quad - \sqrt{\eta_4} (\sqrt{\eta_5} \sqrt{1 - \eta_3} t_H - \sqrt{\eta_5} \sqrt{\eta_3} t_V + \sqrt{1 - \eta_5} v_t) \end{array} \right.$$

To solve the system of coupled equations and to implement a CNOT transformation whose success probability is independent of the input state, it is necessary to take into account the following requirements.

- When the control is $|H\rangle_c$, the mode c_V is empty and must not mix with t_H or t_V .
- When the control is $|V\rangle_c$, the mode c_V must mix symmetrically with both t_H and t_V to produce the swap $t_H \leftrightarrow t_V$, namely the CNOT flip on the target.
- All four postselected amplitudes must have equal modulus, so that the gate acts unitarily on the postselected subspace regardless of the input state.
- By construction, the gate treats the two target states t_H and t_V symmetrically; therefore, the swap $t_H \leftrightarrow t_V$ must occur with the same amplitude for both.

This last condition leads to $\eta_3 = \eta_4 = \eta'$ and to inspect the Heisenberg equation for c_{Vo} because the coefficient of t_H and of t_V must be equal.

Therefore the system of equations reduces to:

$$\left\{ \begin{array}{l} c_{Ho} = \sqrt{1 - \eta_1} v_c + \sqrt{\eta_1} c_H, \\ c_{Vo} = -\sqrt{\eta_2} c_V + \sqrt{1 - \eta_2} (\sqrt{\eta'} t_H + \sqrt{1 - \eta'} t_V), \\ \sqrt{\eta'} \sqrt{1 - \eta_2} = \sqrt{1 - \eta_2} \sqrt{1 - \eta'}, \\ t_{Ho} = \sqrt{\eta'} (\sqrt{1 - \eta_2} c_V + \sqrt{\eta_2} \sqrt{\eta'} t_H + \sqrt{\eta_2} \sqrt{1 - \eta'} t_V) + \\ \quad + \sqrt{1 - \eta'} (\sqrt{\eta_5} \sqrt{1 - \eta'} t_H - \sqrt{\eta_5} \sqrt{\eta'} t_V + \sqrt{1 - \eta_5} v_t), \\ t_{Vo} = \sqrt{1 - \eta'} (\sqrt{1 - \eta_2} c_V + \sqrt{\eta_2} \sqrt{\eta'} t_H + \sqrt{\eta_2} \sqrt{1 - \eta'} t_V) - \\ \quad - \sqrt{\eta'} (\sqrt{\eta_5} \sqrt{1 - \eta'} t_H - \sqrt{\eta_5} \sqrt{\eta'} t_V + \sqrt{1 - \eta_5} v_t) \end{array} \right.$$

Then, by postselection we only retain the terms in which exactly one photon is found in each logical output port, with the vacuum modes v_c and v_t that are left empty. In other words, for each state of the computational basis, given by the states $\{|HH\rangle, |HV\rangle, |VH\rangle, |VV\rangle\}$, we take the combination of the output equations, expand the calculations and keep only the amplitude of the combination of the desired input modes. Inspecting these Heisenberg equations, the four postselected amplitudes are obtained as follows:

$$\begin{aligned}
|HH\rangle &\rightarrow |HH\rangle : \quad \eta' \sqrt{\eta_1 \eta_2} + (1 - \eta') \sqrt{\eta_1 \eta_5}, \\
|HV\rangle &\rightarrow |HV\rangle : \quad \eta' \sqrt{\eta_1 \eta_5} + (1 - \eta') \sqrt{\eta_1 \eta_2}, \\
|VH\rangle &\rightarrow |VV\rangle : \quad -\sqrt{\eta_2}(\sqrt{\eta_2 \eta'(1 - \eta')} - \sqrt{(1 - \eta') \eta' \eta_5}) + (1 - \eta_2) \sqrt{\eta'(1 - \eta')} \\
|VV\rangle &\rightarrow |VH\rangle : \quad -\sqrt{\eta_2}(\sqrt{\eta_2 \eta'(1 - \eta')} - \sqrt{(1 - \eta') \eta' \eta_5}) + (1 - \eta_2) \sqrt{\eta'(1 - \eta')}
\end{aligned}$$

If we want that all these four amplitudes be equal, we have to impose

$$\eta' \sqrt{\eta_1 \eta_2} + (1 - \eta') \sqrt{\eta_1 \eta_5} = \eta' \sqrt{\eta_1 \eta_5} + (1 - \eta') \sqrt{\eta_1 \eta_2} = \sqrt{\eta'(1 - \eta')}(\sqrt{\eta_2 \eta_5} + 1 - 2\eta_2) \quad (2.3)$$

Combined with the condition on c_{Vo} : $\sqrt{\eta'} \sqrt{1 - \eta_2} = \sqrt{1 - \eta_2} \sqrt{1 - \eta'}$, we obtain the following simplified system to solve:

$$\begin{cases} \sqrt{\eta'} \sqrt{1 - \eta_2} = \sqrt{1 - \eta_2} \sqrt{1 - \eta'} \\ \eta' \sqrt{\eta_1 \eta_2} + (1 - \eta') \sqrt{\eta_1 \eta_5} = \eta' \sqrt{\eta_1 \eta_5} + (1 - \eta') \sqrt{\eta_1 \eta_2} \\ \eta' \sqrt{\eta_1 \eta_5} + (1 - \eta') \sqrt{\eta_1 \eta_2} = \sqrt{\eta'(1 - \eta')}(\sqrt{\eta_2 \eta_5} + 1 - 2\eta_2) \end{cases} \quad (2.4)$$

From the first equation, we immediately obtain that $\eta' = 1/2$, which allows to straightforwardly simplify it from the other equations of the system, thus obtaining:

$$\begin{cases} \eta' = \frac{1}{2} \\ \sqrt{\eta_1 \eta_2} + \sqrt{\eta_1 \eta_5} = \sqrt{\eta_1 \eta_5} + \sqrt{\eta_1 \eta_2} \\ \sqrt{\eta_1 \eta_5} + \sqrt{\eta_1 \eta_2} = \sqrt{\eta_2 \eta_5} + 1 - 2\eta_2 \end{cases} \quad (2.5)$$

The second equation is an identity, which means that the original system of equations does not possess a unique solution.

Indeed, in Sec. III of their paper, Ralph et al. write: *"For simplicity we assume that the beam splitters all came from the same 'production run' such that any deviation from the optimal value is common[...]"* [10]. Thus, they assume that B_1 , B_2 , and B_5 share the same actual reflectivity value (η), rather than each carrying an independent error. This choice is justified by the fact that beam splitters from the same production run tend to exhibit systematic, correlated deviations from the nominal value rather than uncorrelated errors, thereby reducing the model free parameters from five to two (η , η').

Hence, with this choice of $\eta = 1/3, \eta' = 1/2$, which clearly satisfies the system in Eq. 2.5, they finally obtain Heisenberg equations of motion as:

$$\begin{cases} c_{Ho} = \frac{1}{\sqrt{3}} (\sqrt{2} v_c + c_H), \\ c_{Vo} = \frac{1}{\sqrt{3}} (-c_V + t_H + t_V), \\ t_{Ho} = \frac{1}{\sqrt{3}} (c_V + t_H + v_t), \\ t_{Vo} = \frac{1}{\sqrt{3}} (c_V + t_V - v_t), \end{cases}$$

Calculation of the CNOT success probability Consider the general two-photon input state

$$|\psi\rangle_{in} = \left(\alpha_1 c_H^\dagger t_H^\dagger + \alpha_2 c_H^\dagger t_V^\dagger + \alpha_3 c_V^\dagger t_H^\dagger + \alpha_4 c_V^\dagger t_V^\dagger \right) |0\rangle,$$

with $\sum_i |\alpha_i|^2 = 1$. Substituting the output operators from the Heisenberg equations and expanding, the output state is a superposition of many terms corresponding to all possible distributions of the two photons across the six output modes. The postselection condition — exactly one photon in each of the two logical output ports, with the vacuum modes v_c, v_t empty — selects one term from each of the four input amplitudes.

For the term $\alpha_1 c_H^\dagger t_H^\dagger$, substituting the output operators gives

$$c_{Ho}^\dagger t_{Ho}^\dagger = \frac{1}{3} \left(\sqrt{2} v_c^\dagger + c_H^\dagger \right) \left(c_V^\dagger + t_H^\dagger + v_t^\dagger \right),$$

of which the only postselected term is $\frac{1}{3} c_H^\dagger t_H^\dagger |0\rangle = \frac{1}{3} |HH\rangle$.

For $\alpha_2 c_H^\dagger t_V^\dagger$:

$$c_{Ho}^\dagger t_{Vo}^\dagger = \frac{1}{3} \left(\sqrt{2} v_c^\dagger + c_H^\dagger \right) \left(c_V^\dagger + t_V^\dagger - v_t^\dagger \right),$$

giving postselected term $\frac{1}{3} |HV\rangle$.

For $\alpha_3 c_V^\dagger t_H^\dagger$:

$$c_{Vo}^\dagger t_{Ho}^\dagger = \frac{1}{3} \left(-c_V^\dagger + t_H^\dagger + t_V^\dagger \right) \left(c_V^\dagger + t_H^\dagger + v_t^\dagger \right).$$

Expanding and retaining only terms with one photon in c_V and one in t_H or t_V (all other modes empty), the terms $-c_V^\dagger t_H^\dagger$ and $+t_H^\dagger c_V^\dagger$ vanish (since $c_V^\dagger$ and $t_H^\dagger$ act on different modes, they commute), leaving $\frac{1}{3} c_V^\dagger t_V^\dagger |0\rangle = \frac{1}{3} |VV\rangle$.

Similarly for $\alpha_4 c_V^\dagger t_V^\dagger$:

$$c_{Vo}^\dagger t_{Vo}^\dagger = \frac{1}{3} \left(-c_V^\dagger + t_H^\dagger + t_V^\dagger \right) \left(c_V^\dagger + t_V^\dagger - v_t^\dagger \right),$$

giving the postselected term $(1/3) c_V^\dagger t_H^\dagger |0\rangle = (1/3) |VH\rangle$.

The full postselected output state is therefore

$$|\psi\rangle_{out.cb} = \frac{1}{3} (\alpha_1 |HH\rangle + \alpha_2 |HV\rangle + \alpha_3 |VV\rangle + \alpha_4 |VH\rangle),$$

which is precisely the CNOT transformation with control and target qubits preserved. The success probability is the squared norm of this state,

$$p_{succ} = \frac{1}{9} (|\alpha_1|^2 + |\alpha_2|^2 + |\alpha_3|^2 + |\alpha_4|^2) = \frac{1}{9},$$

where the last equality follows from the normalisation of the input state. The value $1/9$ is thus input-state independent, a consequence of the equal-amplitude constraint that fixed $\eta = 1/3$ in the first place. If we had set different values of η_1, η_2, η_5 , we would have had a complicated probability calculation with final value dependent on α_i , so not constant.

We notice that the success probability of $1/9$ is actually lower than the $1/4$ achievable with a two-photon ancilla [9], which reflects the cost of eliminating the ancilla photons altogether. More importantly, this gate is fundamentally not feed-forwardable: since the

coincidence condition that signals success requires detecting the logical output photons themselves, the gate cannot be incorporated into a larger circuit without destroying the qubits it is supposed to process.

In fact, this is a subtle point that deserves a targeted discussion. In fact, this CNOT is unsuitable for scalable quantum computation in the KLM sense due the reason above. However, it may be well suited as a Bell-state analyzer: by adding a 50:50 beam splitter to the control output, and then measuring in the superposition basis, the gate can distinguish all four Bell states, with direct applications in quantum communication protocols such as quantum teleportation and entanglement-based cryptography. We notice that the authors also analyse the sensitivity of the gate to imperfections of the beam splitting elements, concluding that error rates below 1% are realistic with state-of-the art beam splitter technology, while mode-matching errors are the dominant experimental limitation.

2.4.1 Experimental realizations: CNOT without ancilla photons

The scheme proposed by Ralph et al. to achieve a probabilistic CNOT gate without the use of any additional ancilla photons other than the ones required for encoding information into the control and target qubits has been realized in different technological versions. Here we report the main achievements in this sense.

Free photon realization (2003). A first implementation proposal of the CNOT gate without ancilla photons, following exactly the scheme introduced above, was made in 2003 by J. O’Brien et al. [13]. The experimental setup is built around an inherently stable interferometer that maps the conceptual CNOT scheme of Fig. 2.6 directly onto a polarization-encoding architecture. It is worth noting that this entire apparatus is implemented in free space (bulk optics): the photons propagate through air, and the interferometer is built from discrete waveplates and beam splitters mounted on an optical table, rather than being inscribed as waveguides on a monolithic chip.

Pairs of energy-degenerate photons, generated via beam-like spontaneous parametric down-conversion and collected into single-mode fibers, are launched from the left of the apparatus (see Fig. 2.6). Each input beam passes through a half-wave plate (HWP) and quarter-wave plate (QWP), allowing any pure, separable two-qubit input state to be prepared. The horizontal and vertical polarization components of the control and target qubits are split and recombined using polarizing beam splitters (PBSs), producing parallel, displaced output modes; this design is intrinsically stable against the sub-wavelength path-length fluctuations that would otherwise plague such a scheme. The target interferometer is realized by mixing and recombining the logical basis modes while the target qubit remains polarization-encoded, using a HWP set to rotate the polarization by 45° , which acts as a Hadamard gate equally mixing the two polarization modes. The action of all three 1/3 beam splitters required by the conceptual scheme is realized by a single 1/3 HWP: the c_1 and t_0 components, orthogonally polarized, exit the first PBS in the same spatial mode and are unequally mixed at this waveplate, where the necessary two-photon quantum interference occurs because the photons become only partially distinguishable after the second PBS. The orientation of this waveplate is

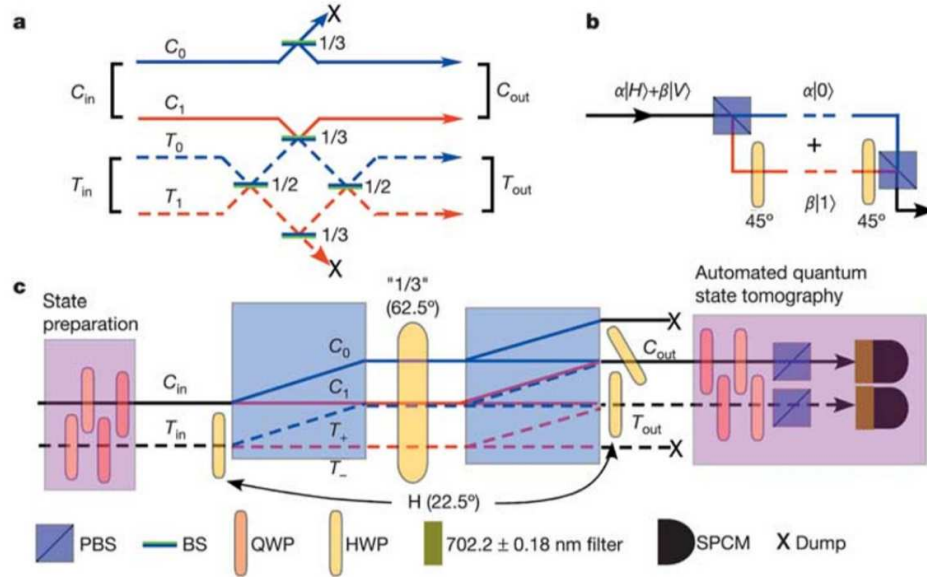
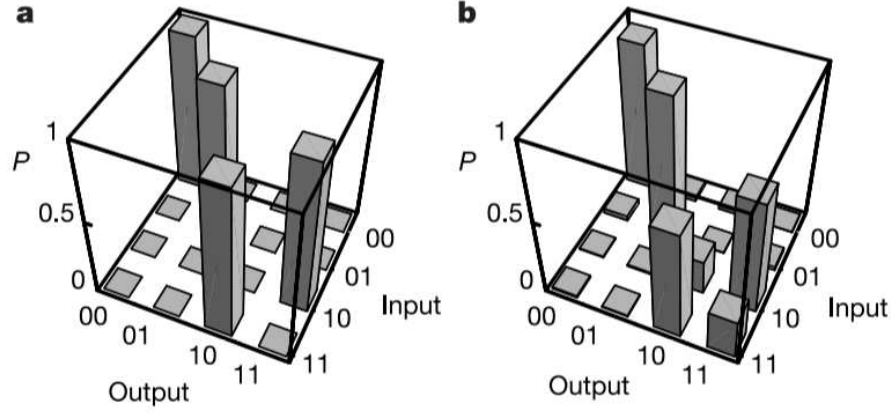


Figure 2.6: A schematic representation of the experimental setup implementing the CNOT gate proposed in Ref. [10]. (a) Conceptual scheme of the gate, as described in the text. A sign change (π -phase shift) occurs upon reflection off the green side of the beam splitters (BSs). (b) A polarization encoded photonic qubit can be converted into a spatially encoded qubit, suitable for the gate shown in (a), using a polarizing beam splitter (PBS) and a half-wave plate (HWP) set to rotate the polarization of one of the outputs by 90° . (c) A schematic of the experimental CNOT gate. Original picture taken from Ref. [13].

chosen both to achieve the required interference and to correctly recombine the control and target modes at the second PBS, while routing $\frac{2}{3}$ of the unwanted components into beam dumps to balance the gate. A fixed HWP is included in the control arm to correct a residual phase shift.

Finally, the two output modes are polarization-analyzed using an automated quantum state tomography system. A coincidence count, simultaneous detection of a single photon at each output, signals successful (heralded) operation of the gate. In the logical basis the gate operates with an average success of 84%, as reported in Fig. 2.7a for completeness. Combined with QND measurement, the CNOT gate demonstrated here is equivalent to the KLM CNOT gate [5] and could, in principle, be rendered scalable through the same teleportation protocol. As such, it actually constituted the first experimental indication that all-optical quantum computation using single-photon qubits could indeed be feasible.

Integrated photonics realization (2008). High-visibility quantum interference requires excellent spatial overlap of the optical modes at a beam splitter, demanding precise mode alignment, while high-visibility classical interference additionally requires sub-wavelength stability of the optical path lengths, often calling for elaborate, carefully stabilized interferometers. Together with photon loss, this interference visibility represents the dominant factor limiting the performance of optical quantum circuits. Hence, a possible technological solution to these problems is integrated photonics. Light is confined within waveguides, each consisting of a core surrounded by a cladding of slightly lower refractive index (analogously to guided electromagnetic signals in optical



(a) Experimental demonstration of CNOT operation in the logical basis. a) Ideal logical basis operation of a CNOT gate; b) Measured operation for the gate presented in Ref. [13].

Table 1 **Experimentally determined probabilities for the logical basis operation**

Input $ CT\rangle$	$P_{ 00\rangle}$	$P_{ 01\rangle}$	$P_{ 10\rangle}$	$P_{ 11\rangle}$
$ 00\rangle$	0.95(2)	0.023(3)	0.024(3)	0.0006(5)
$ 01\rangle$	0.031(3)	0.94(2)	0.0019(8)	0.022(3)
$ 10\rangle$	0.005(1)	0.011(2)	0.23(9)	0.75(2)
$ 11\rangle$	0.011(2)	0.0005(1)	0.72(2)	0.26(1)

(b) Experimental data (corresponding to the ones plotted in the figure above on the right), as taken from the original Ref. [13].

Figure 2.7: All the pictures, data, and table shown in this Figure have been taken from the original reference [13].

fibers). These devices are typically fabricated on semiconductor substrates. Through appropriate choice of the core and cladding dimensions, as well as of their refractive index contrast, such waveguides can be designed to support a single transverse mode over a given wavelength range. Beam-splitter-like coupling between two waveguides is achieved by bringing them close enough together that their evanescent fields overlap, a configuration known as a directional coupler, as shown in Fig. 2.8a(a). The fraction of light transferred from one waveguide to the other, i.e., the coupling ratio $1 - \eta$, with η playing the role of beam-splitter reflectivity, can be precisely tuned by lithographically setting the separation between the waveguides and the length of the coupling region.

A realization of this type of device was made in 2008 by Politi et al. in glass material [14], which explicitly mimics the bulk optics realization previously made by J. O’Brien et al. in Ref. [13]. High-fidelity silica-on-silicon integrated optical realizations of various quantum photonic circuits have been demonstrated in Ref. [14], including a CNOT gate with an average logical basis fidelity of $94.3 \pm 0.2\%$, showing that it is possible to implement sophisticated photonic quantum circuits onto a silicon chip.

The material platform was chosen to satisfy three requirements: *i)* low optical loss at $\lambda \sim 800$ nm, where commercial silicon avalanche photodiode SPCMs are near peak efficiency; *ii)* a refractive index contrast Δ enabling single-mode operation for waveguide dimensions comparable to standard single-mode fiber cores (4–5 μm), for good

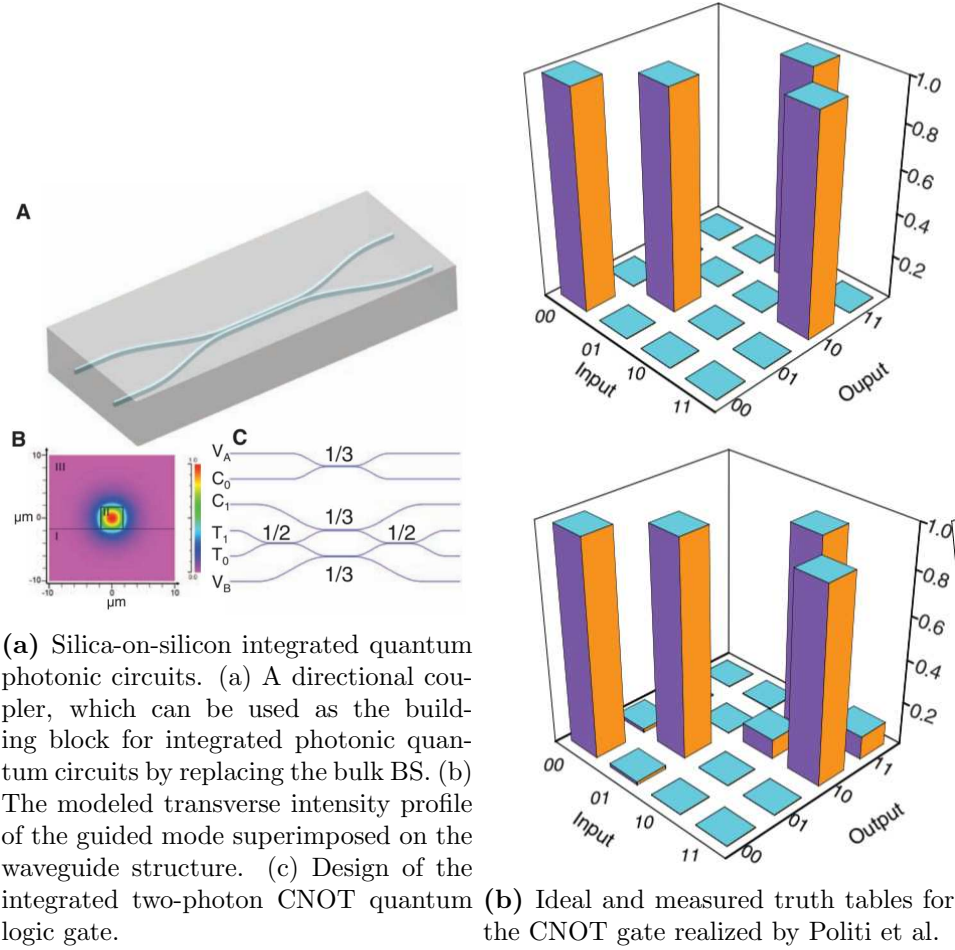


Figure 2.8: All the pictures, data, and table shown in this Figure have been taken from the original reference [14].

fiber coupling; *iii*) and compatibility with standard optical lithography. Silica (SiO_2), doped to control its refractive index and grown on a silicon substrate, best met these criteria. A contrast of $\Delta = 0.5\%$ was selected to give single-mode operation at 804 nm for $3.5 \times 3.5 \mu\text{m}$ waveguides, providing moderate mode confinement that minimizes sensitivity to fabrication and modeling imperfections, as shown in Fig. 2.8a(b).

The chips were fabricated starting from a 4-inch silicon wafer: a $16\text{-}\mu\text{m}$ thermally grown undoped silica buffer layer, followed by flame hydrolysis deposition of a $3.5\text{-}\mu\text{m}$ germanium- and boron-doped silica core, patterned into waveguides by standard lithography, and finally overgrown with a $16\text{-}\mu\text{m}$ phosphorus- and boron-doped silica cladding index-matched to the buffer. The wafer was then diced into individual chips, some of which were polished to improve fiber coupling.

Degenerate photon pairs at 804 nm were generated by type-I spontaneous parametric down-conversion in a beta-barium borate crystal pumped by a 402-nm continuous-wave diode laser, collected into polarization-maintaining fibers, and filtered with 2-nm interference filters to ensure spectral indistinguishability. Photons were launched into the chip and collected via fiber arrays butt-coupled to the waveguides with index-matching fluid, achieving an overall coupling efficiency of $\sim 60\%$.

The integrated CNOT gate implemented on-chip, as shown in Fig. 2.8a(c), consists of

two 1/2 couplers and three 1/3 couplers (the same as in theoretical scheme of Ralph et al. [10]). As in the bulk-optics case [13], successful (heralded) operation occurs with probability 1/9, signaled by detection of a photon at both the control and target outputs. For the nominal device ($\eta_{1/2} = 0.5$, $\eta_{1/3} = 0.33$), the four computational basis states $|0\rangle_C |0\rangle_T$, $|0\rangle_C |1\rangle_T$, $|1\rangle_C |0\rangle_T$, $|1\rangle_C |1\rangle_T$ were input and the output probabilities measured, yielding peak values of 98.5% for the $|0\rangle_C$ inputs and an average logical basis fidelity of $F = 94.3 \pm 0.2\%$ (see Fig. 2.8b).

Several CNOT devices with 1/2 couplers ranging from $\eta_{1/2} = 0.4$ to 0.6 (and correspondingly scaled 1/3 couplers) were also fabricated on the same chip; their fidelities were found to be lower than that of the nominal device, consistent with deviation from the ideal coupling ratios. This high fidelity for the $|0\rangle_C$ inputs reflects the subwavelength phase stability inherent to the monolithic waveguide architecture, in contrast to the bulk-optics implementation, where such stability must be actively engineered.

Summary

This chapter brought together the KLM protocol, the theoretical bounds that followed from it, and the various dedicated constructions of a linear-optical CNOT gate, ending with two experimental realizations of the same ancilla-free scheme originally proposed by Ralph et al. [10]: first, J. O’Brien’s bulk-optics implementation, and then A. Politi’s integrated version on a silicon chip. Comparing all these works makes a point that will come up again: getting high-fidelity, stable interference between photons is at least as much an engineering problem as a theoretical one.

In the next chapter, we will examine two further approaches that have been recently proposed to implement a linear-optical CNOT gate: a deterministic, modular scheme built from cascaded directional-coupler blocks, and a probabilistic scheme based on continuous-time quantum walks. Together, these two case studies illustrate the trade-off between determinism and architectural simplicity, which will be the starting point of our original thesis work.

Chapter 3

Nonlinear quantum walks and interferometers

In this third Chapter, we discuss two recent and conceptually distinct approaches proposing the implementation of two-qubit gates in photonic systems, which constitute the starting point of the original work presented in this thesis.

The first proposal exploits correlated quantum walks in waveguide lattices: two indistinguishable photons propagating through a one-dimensional array of coupled waveguides develop non-classical spatial correlations that can be shaped, through a suitable choice of the lattice parameters, to implement probabilistic quantum logic with a physically compact and simple design [15].

The second approach goes in a different direction: by introducing a distributed self-Kerr nonlinearity in a photonic interferometer, it is shown that fully deterministic two-qubit gates with fidelities arbitrarily close to 100% can be achieved, thus overcoming the fundamental limitations of linear optics for quantum computing applications. However, this result relies on a multi-block concatenated architecture, which raises questions about practical scalability [17].

Hence, the natural question that motivates our work is whether the relative simplicity of the quantum walk architecture and the power of nonlinear interactions can be somehow combined to achieve better performance in a realistically attainable platform.

3.1 Quantum walks and probabilistic two-qubit gate

The concept of a *quantum random walk* was introduced by Aharonov, Davidovich and Zagury in 1993 as the quantum-mechanical counterpart of the classical random walk [64]. As a reminder, we notice that in a classical random walk, a particle steps left or right at each iteration according to fixed probabilities. In its quantum version, this rule is replaced by probability amplitudes: before any measurement is performed, the particle exists in a superposition of both directions, and these amplitudes can interfere. This is made possible by coupling the spatial motion of the particle to an auxiliary two-level system, which plays the role of a *quantum coin*. A single step of the walk is implemented by a unitary operator that displaces the particle to the right if the coin is in state $|\uparrow\rangle$, and to the left if it is in state $|\downarrow\rangle$. As long as the coin is not measured, the two displacements coexist coherently and interfere. In the discrete-time formulation,

the walker occupies one of a discrete set of sites $\{\dots, -1, 0, 1, \dots\}$ on a line, and one step of the walk is the composition of a coin-flip C followed by a conditional shift S . If the coin is measured after every iteration, the process reduces to an ordinary classical random walk. A genuine quantum walk is instead obtained by iterating C and S without any intermediate measurement, allowing amplitudes associated with different paths to interfere coherently. The result is a variance that grows quadratically with the number of steps, $\sigma^2 \sim T^2$, corresponding to a ballistic rather than diffusive propagation (a striking departure from the classical case where $\sigma^2 \sim T$ [65]).

Quantum walks in waveguide lattices The discrete-time quantum walk on a line has a direct physical realization in integrated photonics: a lattice of evanescently coupled waveguides. In this system, each waveguide corresponds to one lattice site, the evanescent coupling between adjacent guides plays the role of the coin-and-shift operation, and the propagation of light along the length of the chip plays the role of time. No active control is required at each step: the entire walk is encoded in the static geometry of the device, and the number of steps is fixed by the physical length of the chip. This equivalence between the quantum walk formalism and light propagation in waveguide lattices was demonstrated experimentally by Perets, Lahini et al. [66], thus establishing that the waveguide lattice is a natural and scalable platform for quantum walk experiments.

In a waveguide lattice, the longitudinal direction (along the chip length) plays the role of time, while the transverse direction (from one waveguide to the adjacent one) plays the role of the discrete spatial coordinate of the walk. As light propagates forward along the chip, the evanescent coupling between neighbouring guides continuously redistributes the amplitude across the lattice, implementing one step of the walk at each incremental propagation length.

However, when a single photon propagates through such a lattice, its statistics can still be reproduced by a classical electromagnetic wave: no genuinely quantum resource is required [67]. The situation changes when two indistinguishable photons are injected into the same lattice. In that case, the joint evolution of the two walkers develops correlations that depend directly on the bosonic statistics of the photons, and these correlations have no classical analogue [67, 68]. This is the photonic analogue of the Hong-Ou-Mandel effect (see Sec. 1.3.3), extended from a single beam splitter to an entire lattice of coupled waveguides evolving over many steps: the result is a rich spatial correlation structure in the joint probability distribution of the two photons at the output of the chip.

It is precisely this two-photon interference in a waveguide lattice that Lahini et al. conceptually exploit to implement a quantum logic gate [15]. By encoding logical qubit states in the positions of two indistinguishable photons propagating through a one-dimensional waveguide array, and by post-selecting on specific coincidence outcomes at the output (the states of the logical base), they show that the natural evolution of the correlated quantum walk implements a probabilistic two-qubit gate.

3.1.1 Probabilistic CNOT gate in linear quantum walks

The starting point of the construction is the Bose-Hubbard Hamiltonian, which governs the dynamics of both physical platforms considered in this work, and can be generally

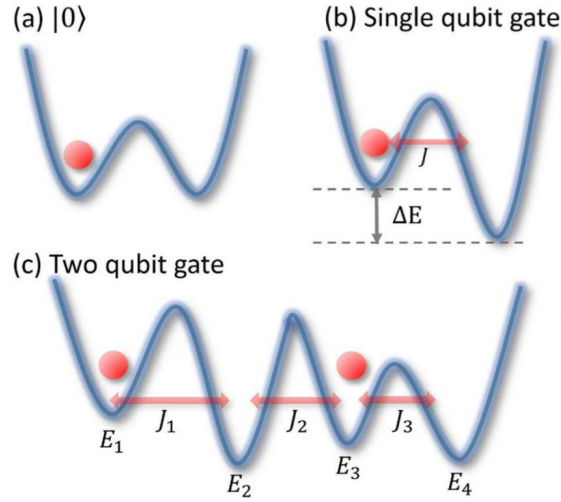


Figure 3.1: Quantum-walk-based quantum logic gates for cold atoms and for photons. In atomic implementations, each well represents an atomic trap potential. **a)** A qubit in the state $|0\rangle$. **b)** Implementation of a single-qubit gate. **c)** Schematic of a two-qubit system on a lattice. Picture taken from the original Ref. [15].

formulated as:

$$\mathcal{H} = \sum_m E_m a_m^\dagger a_m + \sum_m J_m (a_m^\dagger a_{m+1} + a_{m+1}^\dagger a_m) + \frac{\Gamma}{2} \sum_m a_m^\dagger a_m (a_m^\dagger a_m - 1), \quad (3.1)$$

in which

- **in the atomic case:** E_m is the on-site energy of the m -site; $a_m^\dagger, a_m$ are the creation/annihilation operators for an atom; $J_m \leq 0$ is the tunnelling rate between the nearest-neighbour sites; Γ is the on-site interaction energy, this value is usually nonzero and can be adjusted experimentally.
- **in the photonic case:** E_m is the on-site energy determined by the index of refraction of the m -site (one-dimensional waveguide); $a_m^\dagger, a_m$ are the creation/annihilation operators for a single photon quantum in the m -th waveguide; $J_m \leq 0$ is the tunnelling rate between the nearest sites; Γ is the on-site photon-photon interaction energy (in the linear photonic case, $\Gamma = 0$).

The continuous-time QW is described by the unitary evolution under the Hamiltonian described by $U_{atomic} = e^{-iHt}$ for the atomic system, and by $U_{photonic} = e^{-iHz}$ for the photonic system.

To define the qubits on the lattice, the usual dual-rail encoding is employed, as already introduced in Chapter 1: a qubit is physically implemented by a single boson in a pair of neighbouring potential wells (see, e.g., Fig. 3.1), with the logical basis states of the single qubit, $|0\rangle$ and $|1\rangle$, defined by the particle being in the left or right well, respectively. With this implementation, a system of n qubits can be realized in one dimension using n bosons and $N = 2n$ lattice sites.

The central challenge is an *inverse problem*: while it is straightforward to compute the unitary evolution operator $U = e^{-iHt}$ from a given set of lattice parameters, the

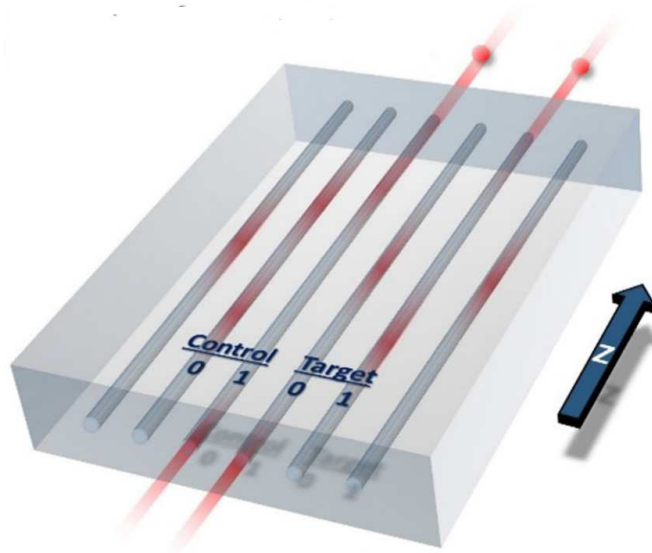


Figure 3.2: Schematic representation of a quantum-walk-based photonic CNOT gate. Picture taken from the original Ref. [15].

converse calculation — i.e., finding the Hamiltonian H that realizes a desired logical gate — is far from trivial. There is no analytic inverse to the matrix exponential, the physical constraints of the device (one-dimensionality, nearest-neighbor coupling, real tunnelling coefficients) severely restrict the space of admissible solutions, and the logical states span only a subspace of the full Hilbert space, leaving U not even uniquely defined by the target gate. Lahini et al. have addressed this difficulty through a combined analytical and numerical approach, optimizing the lattice parameters via nonlinear algorithms to maximize the gate fidelity over the universal set given by {single-qubit rotations, CNOT}.

In this approach, single-qubit gates are implemented analytically: since they involve only one particle in two sites, the interaction Γ is irrelevant and the lattice parameter matrix $H_{\text{single-particle}}$ coincides directly with the single-particle Hamiltonian.

$$H_{\text{single-particle}} = \begin{pmatrix} E_1 & J_{12} \\ J_{12} & E_2 \end{pmatrix}$$

The desired gate is then obtained as $U = e^{-itH_{\text{single-particle}}}$, with the on-site energies and tunneling coefficient tuned to reproduce the target transformation.

As we have already mentioned, a universal gate set for quantum computation includes the CNOT, together with all single-qubit unitaries (see, e.g., Chapter 1). For the two-qubit gates, the authors found the lattice parameters via numerical optimization, by using a combination of a stochastic global search algorithm and a gradient-free local optimizer. The gate fidelity is defined through the Hilbert-Schmidt inner product between the target unitary, U_0 (i.e., the theoretical CNOT gate) and the optimized unitary U , restricted to the logical subspace of dimension $N = 4$:

$$F(U_0, U) = |\langle U_0, U \rangle|, \quad \langle U_0, U \rangle = \frac{\text{Tr}(U_0^\dagger U)}{N}, \quad (3.2)$$

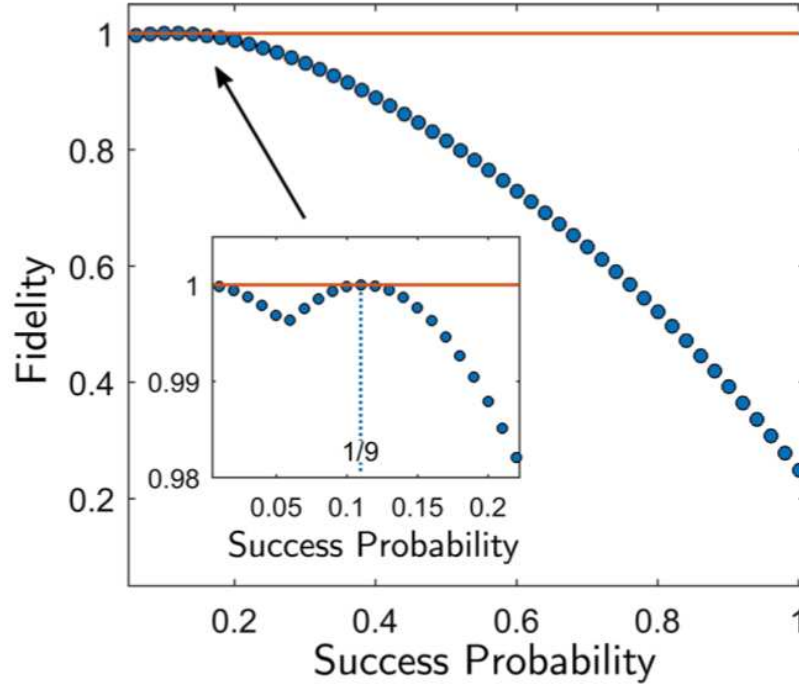


Figure 3.3: Fidelity of the photonic CNOT gate as a function of success probability. The Fidelity peaks at 1 for success probability of $1/9$. Picture from [15].

which can be interpreted as a lower bound on the average gate fidelity. The cost function minimized at each optimization step is $1 - F^2$ rather than $1 - F$, which was found to yield superior convergence, as the local optimizer assumes a quadratic cost function. An additional penalty term $\sin(\arg(u_{1,1}))^2$ is included to fix the global phase of U to match that of U_0 , without affecting the fidelity.

CNOT gate by optimizing non-interacting QWs.

The non-interacting case ($\Gamma = 0$) corresponds to indistinguishable photons propagating in a waveguide lattice, where no effective photon-photon interaction is available. As discussed in Chapter 1, the absence of interactions forces any linear-optical gate to be probabilistic: the device can only implement the desired logical transformation on a subset of the output events, identified by post-selection on specific coincidence outcomes. The authors therefore allow for additional auxiliary (vacuum) sites beyond the four sites required to encode two qubits, extending the lattice to n sites with $n - 1$ coupling coefficients, and optimize simultaneously over gate fidelity and success probability. Given the result of Ralph et al. [10], the starting configuration to be optimized by Lahini et al. is a six-waveguide device, in which only the central four devices constitute the computational basis states, and are initialized with two single-photons at the input (see Fig. 3.2 for a schematic illustration).

The key result from the optimization of the linear QW is shown in Fig. 3.3: a 100% fidelity CNOT gate exists at a success probability of exactly $1/9$, saturating the bound derived by Ralph et al. [10] for linear-optical CNOT gates in the coincidence basis. This

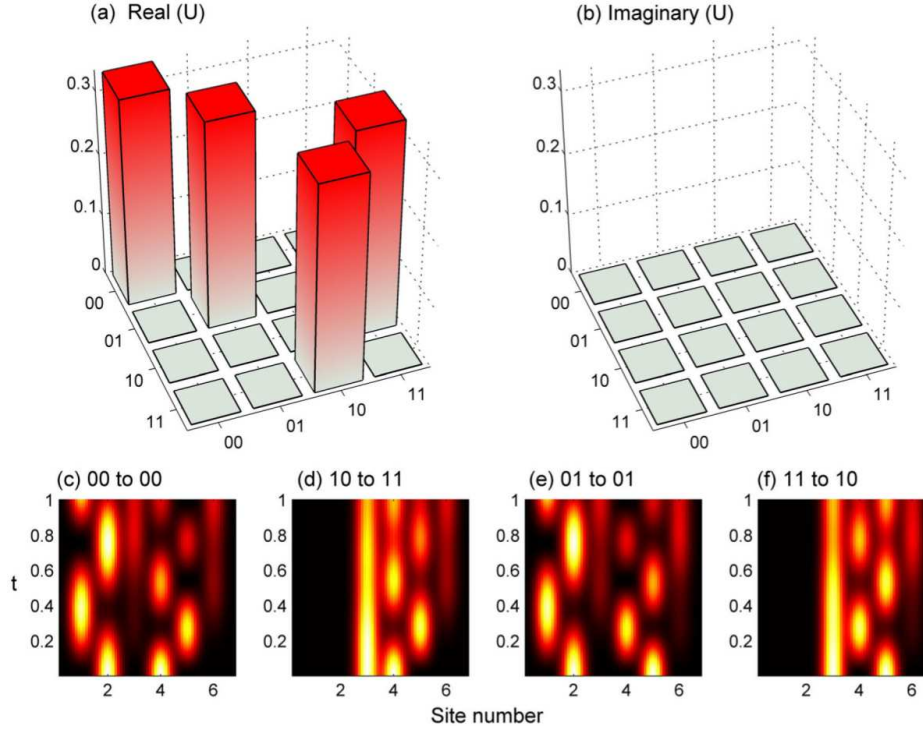


Figure 3.4: A quantum-walk-based photonic CNOT gate. **a), b)** The real and imaginary values of the relevant sub-matrix of the two-particle unitary generated by the Hamiltonian in Eq.3.1, yielding a photonic CNOT gate with success probability of $1/9$. Plots **c)-f)** show the photon amplitudes as a function of time. Picture from [15].

is achieved on a six-site lattice with optimal Hamiltonian parameters obtained as:

$$H_{\text{opt}} = \pi \begin{pmatrix} 0 & -1.27 & 0 & 0 & 0 & 0 \\ -1.27 & -0.73 & 0 & 0 & 0 & 0 \\ 0 & 0 & -0.67 & -0.51 & 0 & 0 \\ 0 & 0 & -0.51 & 0.01 & -1.69 & 0 \\ 0 & 0 & 0 & -1.69 & -1.01 & -0.52 \\ 0 & 0 & 0 & 0 & -0.52 & -1.67 \end{pmatrix}, \quad (3.3)$$

where the diagonal entries represent the on-site (dimensionless) energies, E_m , the off-diagonal entries are the nearest-neighbour (dimensionless) tunnelling coefficients $J_{m,l}$, and sites 1 and 6 are the two auxiliary vacuum sites.

Crucially, all waveguides are straight and parallel one-dimensional channels (see Fig. 3.2), requiring no complicated designs for tapers and bends, which is in contrast to the beam splitters-based architecture of previous implementations (e.g., Ref. [14]). In fact, bent waveguides may introduce additional propagation losses, which is detrimental for the overall fidelity of the gates, of course. Interestingly, the QW architecture thus optimized reaches the same fundamental bound predicted by Ralph et al., but in a significantly simpler and more compact physical design.

Figures 3.4(a,b) show the real and imaginary parts of the relevant 4×4 sub-matrix of the whole unitary evolution matrix, U , confirming the correct CNOT structure on the logical basis, while the bottom panels in Fig. 3.4(c-f) display the photon amplitude evolution along the one-dimensional QW for each logical input state. As an additional

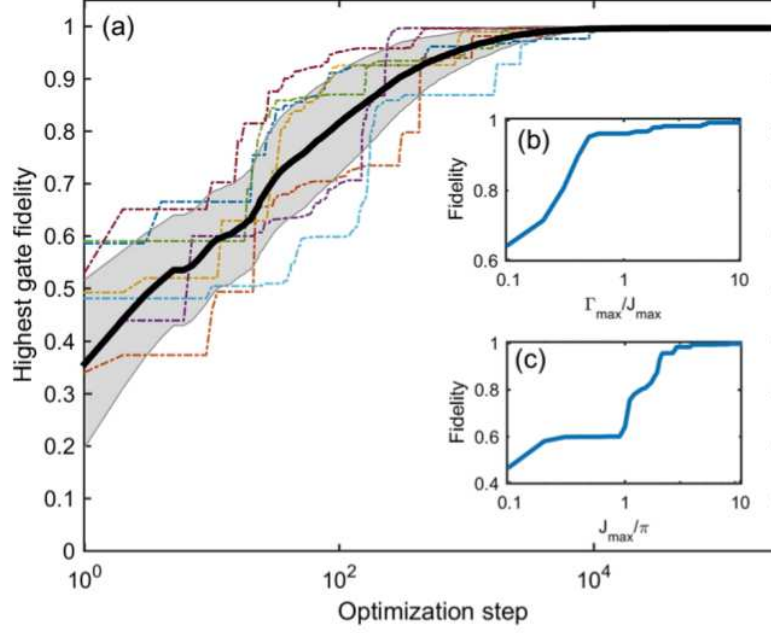


Figure 3.5: Optimization of the deterministic CNOT gate fidelity. **a)** Convergence of different optimization runs to the optimal gate fidelity. The solid black line and the shaded area represents the average and the standard deviation values over 512 runs. Seven example runs are shown in the background (dotted lines). **b)** Gate fidelity versus the maximum allowed interaction level Γ_{max} , at a constant $J_{max} = 4\pi$. **c)** Gate fidelity for different maximal tunneling rates J_{max} at a constant maximal interaction level of $\Gamma_{max} = 20\pi$. Original picture taken from Ref. [15].

feature emerging from Fig 3.3, we notice that the fidelity remains close to unity up to success probabilities of approximately 2/9, suggesting a possible trade-off between single-gate fidelity and heralding efficiency that could be exploited in practical implementations.

Deterministic CNOT gate using strongly interacting QWs.

For the interacting case (i.e., $\Gamma \neq 0$), which in practical implementations may correspond to ultra-cold bosonic atoms in an optical lattice, the CNOT gate is realized on a *four-site lattice* with two bosons, whose full Hilbert space has dimension 10. In this case, the system is described by eight free parameters: four on-site energies E_m , three nearest-neighbor tunneling coefficients $J_{l,m}$, and the on-site interaction strength Γ , optimized numerically under physically motivated bounds:

$$J_{l,m} \in [-J_{max}, 0] \quad |E_m| \leq J_{max} \quad \Gamma \leq \Gamma_{max}$$

with $J_{max} = 4\pi$ limiting the maximum number of tunnelling events during the gate operation. The CNOT acts as a 4×4 sub-matrix of the full evolution operator $U = e^{-iH}$, restricted to the logical states $\{|1010\rangle, |1001\rangle, |0110\rangle, |0101\rangle\}$ in occupation-number basis.

The optimization reported from the authors robustly converges across independent runs to a gate fidelity of 99.6%, with interaction strength $\Gamma = 21.68\pi$ (which is a really high

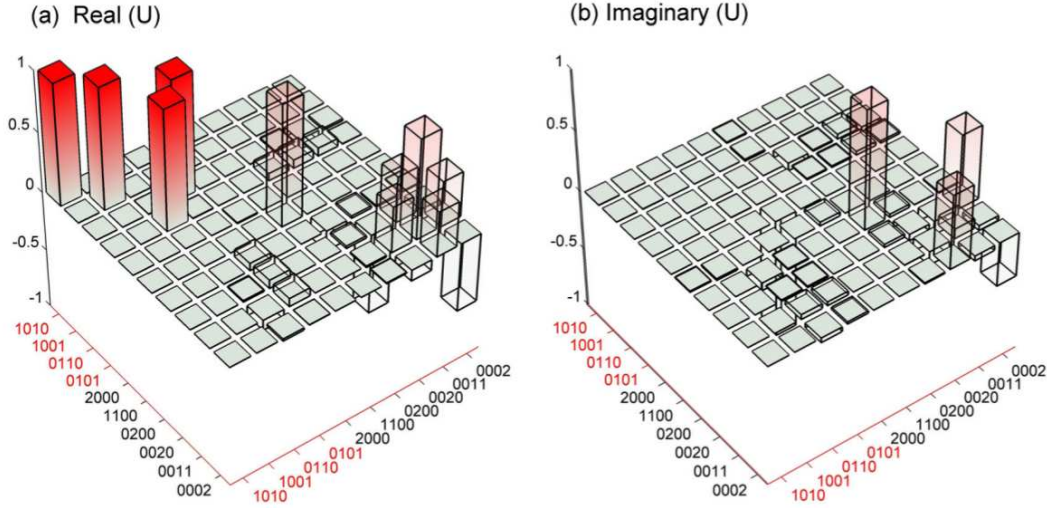


Figure 3.6: **a)** The real part and **b)** the imaginary part of the two-particle unitary transform, U . The CNOT gate operation corresponds to the sub-matrix of the logic states, shown in solid-color bars and marked with red axis labels. Picture from [15].

value, unrealistic in photonic implementations). Although, Fig. 3.5 shows that the fidelity approaches unity as $\Gamma_{\max}$ and $J_{\max}$ are relaxed, confirming that the residual infidelity is a consequence of the experimental bounds rather than a fundamental limitation of the architecture. The optimal CNOT gate is shown in Fig. 3.6, which clearly shows that the gate is deterministic within the four-dimensional space spanned by the computational basis states.

Remarkably, the results from Lahini et al. demonstrate that one-dimensional quantum walks are a surprisingly powerful substrate for quantum logic: a minimal physical resource suffices to realize a complete universal gate set. In the photonic case, the waveguide architecture achieves the fundamental $1/9$ success probability bound with a design that is significantly simpler and more compact than previous beam splitter-based implementations. In the strongly interacting case, the on-site interaction Γ unlocks deterministic operation, yielding fidelities above 99% with experimentally accessible parameters in atomic platforms, but rather unrealistic if one has integrated photonic implementations in mind.

On the other hand, the photonic gates remain intrinsically probabilistic, inheriting all the scalability challenges of linear-optical quantum computing discussed in Chapter 2. The deterministic interacting case, while conceptually powerful, requires ultra-cold atomic systems with single-site resolution and tunable interactions, which are experimentally demanding.

3.1.2 Experimental realization with non-interacting QWs.

It is very interesting to highlight a recent work from the literature, in which the theoretical proposal from Lahini et al. has been experimentally realized in an integrated photonic platform. In particular, Chapman et al. [16] have fabricated six-waveguide QW array in a lithium niobate-on-insulator (LNOI) material platform, exactly fulfilling the optimal CNOT parameters found in Ref. [15]. The choice for LNOI is motivated

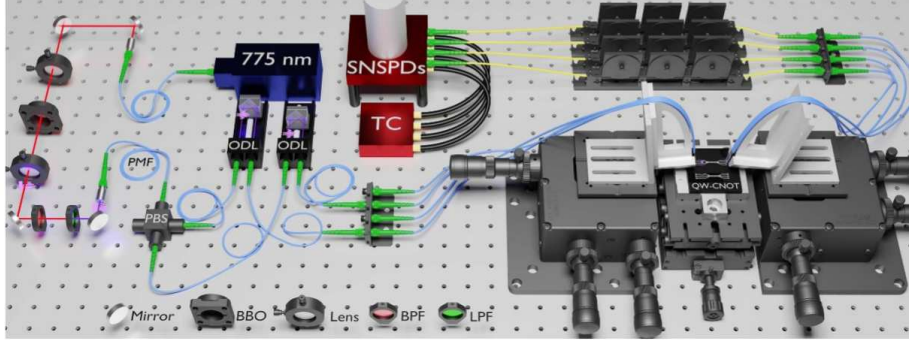


Figure 3.7: Experimental setup from Chapman et al.: Orthogonally polarized photon pairs are generated by type-2 spontaneous parametric down-conversion in a 2 mm thick beta barium borate (BBO) crystal and the pump is filtered with a longpass filter (LPF) and the photons with a 12 nm bandpass filter (BPF). The photons are coupled to polarization maintaining fiber (PMF) and separated with a polarizing beamsplitter (PBS) where one fiber pigtail is rotated such that both photons have the same output polarization. The photons are injected into optical delay lines (ODLs) to compensate for fiber length mismatch and fiber birefringence, and are coupled to the QW-CNOT chip with a fiber array positioned above on-chip grating couplers. After the QW, the photons are collected in fiber and measured with superconducting nanowire single photon detectors (SNSPDs). The arrival times are recorded with a time correlator (TC). Polarization controllers are used to maximize the efficiency of the SNSPDs. Original picture of the experimental table is taken from Ref. [16].

by its broad transparency range, low propagation losses, and large electro-optic and $\chi^{(2)}$ nonlinear coefficients, making it a leading technology for integrated quantum photonics. The tight-binding Hamiltonian of Eq. 3.1 is engineered directly in the waveguide geometry: the on-site energies, $E_i = \beta_i = 2\pi n_{\text{eff},i}/\lambda$, are controlled via the individual waveguide widths, while the nearest-neighbour coupling rates, $J_{ij} = \kappa_i$, are due to evanescent coupling and set by the waveguide separations. The target Hamiltonian-length product, HL , is matched to the matrix in Eq. 3.3 over a total array length of $L = 700 \mu\text{m}$, yielding a device footprint of $700 \mu\text{m} \times 65 \mu\text{m}$.

Indistinguishable photon pairs at 1550 nm are generated via type-2 spontaneous parametric down-conversion (SPDC) in a beta-barium borate crystal. The experimental setup is shown in Fig. 3.7: photons are separated by polarization, routed through motorized optical delay lines to match their arrival times, and injected into the chip via grating couplers. Coincidence measurements are performed with superconducting nanowire single-photon detectors. The indistinguishability of the photon source is verified through a Hong-Ou-Mandel dip with 94.6% visibility.

The central result of this work is the measurement of the two-qubit transfer matrix, shown in Fig. 3.8. When the two photons arrive with a time delay $\tau \neq 0$, no quantum interference occurs and the device behaves classically, with a transfer matrix fidelity of 0.994 ± 0.001 relative to the theoretical model. When the photons arrive simultaneously ($\tau = 0$), quantum interference is activated and the CNOT operation is implemented, yielding a two-qubit transfer matrix fidelity of 0.938 ± 0.003 . The gap between these two values is attributed to the finite bandwidth of the SPDC source (12 nm), which introduces chromatic dispersion in the fibers and waveguides and reduces the visibility of two-photon interference. The fidelity between the simulation — constructed from the

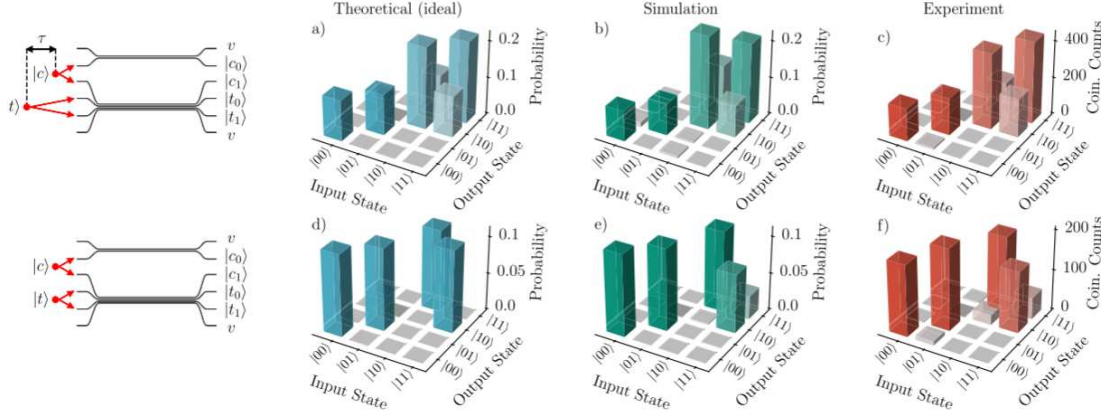


Figure 3.8: Two-qubit transfer matrix measured from Chapman et al. (a)-(c) The theoretical (ideal), simulated and experimentally measured two-qubit transfer matrix when the photons arrive with a time separation τ . In this regime, the photons do not quantum interfere. (d)-(f) The theoretical (ideal), simulated and experimental two-qubit transfer matrix with indistinguishable photons that arrive with zero time delay. Picture taken from Ref. [16].

classically characterized Hamiltonian — and the experiment is 0.985 ± 0.002 , confirming that the main source of infidelity is the photon source rather than the device itself.

It is worth highlighting that these results constitute the first experimental demonstration of a quantum logic gate realized in a continuous-time photonic quantum walk, confirming the theoretical predictions of Lahini et al. in an actual integrated device. The key achievement is that the compact multi-mode interaction region of the waveguide array replaces the network of five individual beamsplitters required by the traditional linear-optical CNOT architecture, with no on-chip routing and no bent waveguides. The measured two-qubit transfer matrix fidelity of 0.938 ± 0.003 demonstrates that the QW architecture is a viable and compact platform for photonic quantum logic. At the same time, the gate remains intrinsically probabilistic, operating with a success probability of $1/9$: a fundamental bound of linear optics that no engineering improvement can overcome. This limitation motivates the search for architectures that go beyond linear optics, introducing nonlinear interactions to achieve deterministic two-qubit gates.

3.2 Quantum nonlinear interferometers and deterministic two-qubit gates

Before turning our discussion to the papers that are central to this Chapter [69, 17], it is essential to first define the *self-Kerr nonlinearity* [70, 71], and understand why it is relevant in photonics.

In conventional classical electro-optical processes, the polarization, $\mathbf{P}$ of a dielectric medium induced by an applied electric field, $\mathbf{E}$, is linearly proportional to its strength, i.e., $\mathbf{P} = \epsilon_0 \chi \mathbf{E}$, where ϵ_0 is the vacuum permittivity and $\chi \equiv \chi^{(1)}$ is the linear susceptibility of the medium, which can be treated as a scalar for linear, homogeneous, and isotropic dielectric media.

For a sufficiently strong field, this linear relation no longer holds, and the induced polarization must be generalized to

$$\mathbf{P} = \epsilon_0 (\chi^{(1)} \mathbf{E} + \chi^{(2)} \mathbf{E} \cdot \mathbf{E} + \chi^{(3)} |\mathbf{E}|^2 \mathbf{E} + \dots) , \quad (3.4)$$

in which $\chi^{(2)}$ and $\chi^{(3)}$ are the second- and third-order nonlinear susceptibilities, respectively (still in scalar approximation, also for the higher order susceptibility contributions). As is typical for nonlinear-optics effects, a clear manifestation of these higher-order terms generally requires high light intensities.

The *Kerr effect* arises from the third-order term in Eq. (3.4), which is the dominant nonlinear effect when the $\chi^{(2)}$ susceptibility can be neglected (such as, e.g., in centrosymmetric crystals): the field experiences an intensity-dependent phase shift, so that the medium behaves as having a refractive index proportional to the field intensity $|\mathbf{E}|^2$, i.e., $n = n_0 + n_2 I$, as already introduced in Chapter 1, Sec. 1.3.2. The coupling χ appearing in the quantum-mechanical treatment below is proportional to this same third-order susceptibility, and hence to n_2 . The self-Kerr nonlinearity is the special case of this effect, in which a single mode interacts with itself.

Unlike many nonlinear optical phenomena, the self-Kerr, cross-Kerr, and Pockels effects do not change the number of excitations in a mode. Instead, they modify the resonance frequency of a mode a . In the self-Kerr effect, the frequency shift is produced by the field in the same mode. In the Pockels effect, the frequency shift is proportional to the amplitude of a second field, $(b + b^\dagger)$. In the cross-Kerr effect, the frequency shift is instead proportional to the photon number of a second mode, $b^\dagger b$. All these effects can be described at the Hamiltonian level by quantizing the electromagnetic field, and then applying the nonlinear response as in Eq. 3.4, from which the three main contributions arise as

$$\begin{aligned} \mathcal{H}_K &= \chi_K (a^\dagger a)^2, \\ \mathcal{H}_{cK} &= \chi_{cK} a^\dagger a b^\dagger b, \\ \mathcal{H}_P &= \chi_P a^\dagger a (b + b^\dagger), \end{aligned} \quad (3.5)$$

which describe the self-Kerr, the cross-Kerr, and the Pockels interactions, respectively, and where χ_i gives the strength of the corresponding nonlinear interaction energy. As discussed in Chapter 1 (see, in particular, Sec. 1.3.2), it is precisely $\mathcal{H}_{cK}$ that underlies the Kerr-mediated CZ gate at the core of the original KLM proposal.

Self-Kerr effect. Focusing on the self-Kerr term, for a single bosonic mode with $[a, a^\dagger] = 1$, the unitary generated by $\mathcal{H}_K$ over a duration t is

$$U(t) = e^{i\chi t a^\dagger a^2 a},$$

where we identify $\chi \equiv \chi_K$. It is convenient to define the dimensionless parameter $\theta = \chi t$, which will be used throughout the rest of this section. Restricting to the two-photon level, the unitary acts on the Fock basis as

$$\alpha |0\rangle + \beta |1\rangle + \gamma |2\rangle \longrightarrow \alpha |0\rangle + \beta |1\rangle + \gamma e^{2i\theta} |2\rangle,$$

where 2θ is the phase shift acquired by the two-photon component.

Single-photon states are left unaffected, while the two-photon component accumulates a phase proportional to the nonlinearity strength and the propagation time. It is precisely this selective phase shift on two-photon states that constitutes the key resource exploited by Scala et al. [17] to generate entanglement between photonic qubits.

3.2.1 Deterministic entangling gates in Kerr-nonlinear interferometers

The theoretical framework underlying the work of Scala et al. [17] was introduced in a preceding paper by Nigro et al. [69], which first analyzed integrated circuits of propagating exciton-polaritons, as a platform for nonlinear quantum photonic gates. Polaritons are hybrid light-matter quasiparticles characterized by an intrinsically large Kerr-type nonlinearity and the Hamiltonian model describing their propagation in two co-directional evanescently coupled waveguides reads as

$$H = \sum_k (H_k^a + H_k^b) + J(x) \left[a_k^\dagger b_k + b_k^\dagger a_k \right], \quad (3.6)$$

where $H_k^\sigma = \omega_k \sigma_k^\dagger \sigma_k + \Gamma_k (\sigma_k^\dagger)^2 \sigma_k^2$ describes the single-waveguide Hamiltonian including the two-body interaction energy Γ_k , and $J(x)$ is the space-dependent evanescent coupling between the two channels.

A key insight of the paper [69] is that the two-polariton dynamics can be mapped exactly onto that of a driven two-level atom, with nonlinearity Γ_k acting as a detuning and the hopping J as an effective Rabi frequency. This atom-polariton correspondence reveals that even weak nonlinearities ($\Gamma_k/J \ll 1$) can have a significant effect on the output two-photon correlations, provided the circuit geometry is appropriately designed.

The most direct physical consequence of this nonlinearity is a modification of the Hong-Ou-Mandel (HOM) effect: in a standard linear beamsplitter, two indistinguishable bosons always bunch at the output. In the presence of polariton interactions, however, the on-site repulsion suppresses the probability of two polaritons occupying the same channel, progressively destroying the HOM dip and driving the system towards an anti-bunching regime in the strongly interacting limit $\Gamma_k \gg J$. It is precisely this competition between hopping and nonlinearity, and the nontrivial two-photon phase shifts it generates, that Scala et al. [17] identify as the key physical mechanism for building deterministic entangling gates. However, the encoding assumed in [69] is of the single-rail type, in which the logical states $|0\rangle$ and $|1\rangle$ correspond to the absence or presence of a photon in a single waveguide channel. As pointed out in [17], unavoidable photon bunching severely hinders the application of single-rail encoding for optical quantum computing, motivating the transition to the dual-rail scheme already discussed in the context of Lahini et al. [15].

These considerations directly motivate the work of Scala et al. [17], which builds on the polariton framework of Nigro et al. [69] while moving to a more general and experimentally flexible setting. They propose a quantum computing paradigm based on dual-rail encoded photonic qubits propagating in nonlinear interferometers, without any restriction to a specific material platform. The central result is that a suitable concatenation of a finite number of elementary interferometric blocks (each combining free propagation and nearest-neighbor hopping regions) is sufficient to implement deterministic two-qubit entangling gates with fidelities arbitrarily close to 100%, without any post-selection. This is demonstrated explicitly for the CNOT and the Mølmer-Sørensen gates, both achieving theoretical fidelities of 99.96% with a circuit depth of around 10–20 blocks.

Crucially, the optimal regime is found at weak nonlinearity relative to the hopping, $\Gamma/J_{\max} \approx 0.5$. This is non-trivial: for $\Gamma \gg J_{\max}$, the system enters the photon block-

ade regime, in which the large energy cost of having two photons in the same channel suppresses precisely the two-photon states through which the qubits interact, hindering rather than enhancing the gate. The value $\Gamma/J_{\max} \approx 0.5$ thus represents the optimal balance between generating sufficient entanglement and avoiding blockade, and is compatible with realistic photonic platforms where single-photon nonlinearities are typically small.

The model system.

The key physical ingredients are a *distributed* self-Kerr nonlinearity Γ , active throughout the entire circuit whenever two photons simultaneously occupy the same waveguide channel, and localized nearest-neighbor hopping J between adjacent channels. The general Hamiltonian describing n propagating channels reads as

$$H = \sum_{i=1}^n \left[\omega_i a_i^\dagger a_i + \Gamma_i (a_i^\dagger)^2 a_i^2 \right] + \frac{1}{2} \sum_{\substack{i,j=1 \\ j \neq i}}^n J_{ij}(x) [a_i^\dagger a_j + a_j^\dagger a_i], \quad (3.7)$$

where ω_i is the energy-momentum dispersion in the i -th channel, $\Gamma_i \equiv \Gamma_i(k)$ is the self-Kerr nonlinearity and $J_{ij}(x)$ are the space-dependent hopping terms between nearest-neighbor channels. For the two-qubit gate implementation, a 4-channel circuit ($n = 4$) is considered, with all channels assumed identical, i.e., $\omega_i = \omega$ and $\Gamma_i = \Gamma$ for all i . The dual-rail encoding maps (see Fig. 3.9a) the two logical qubit states onto single-photon Fock states in pairs of adjacent channels:

$$|0\rangle_2 = |0\rangle \equiv |1, 0\rangle, \quad |1\rangle_2 = |1\rangle \equiv |0, 1\rangle,$$

so that a two-qubit state is represented by two single photons propagating in a 4-channel device, and the computational basis is $\mathcal{S} = \{|00\rangle, |01\rangle, |10\rangle, |11\rangle\}$.

A single-qubit operation emerges naturally from the hopping between two adjacent channels. In the single-photon subspace, the evanescent coupling between the two channels of a qubit implements a rotation around the x -axis of the Bloch sphere, R_X (see eq.(1.15)), with rotation angle $2\theta = \Theta = 2Jt$:

$$R_x(-2\theta) = R_x(-\Theta) = (\cos(\Theta/2)\mathbb{I} + i \sin(\Theta/2)\sigma_X) \begin{bmatrix} a_1^\dagger \\ a_2^\dagger \end{bmatrix},$$

where t is the time spent in the hopping region (see Fig. 3.9b). Combined with a phase shifter that implements R_Z (see eq.(1.16)), this gives a complete set of single-qubit operations. The linear part of the circuit is therefore already sufficient for arbitrary single-qubit rotations; the nonlinearity is only needed for the two-qubit entangling operation.

For the two-qubit gate, a 4-channel circuit is considered, where channels 1–2 encode the first qubit and channels 3–4 encode the second. The circuit is built by concatenating M identical elementary blocks. Each block consist of five types of unitaries acting on the 4-channel device, as sketched in Fig. 3.9c. In the *hopping regions* (HR), two adjacent channels are evanescently coupled with a tunable coupling rate J , implementing a beam-splitter-like interaction between neighboring modes

$$U_{HR}(t) = R_X(2Jt), \quad (3.8)$$

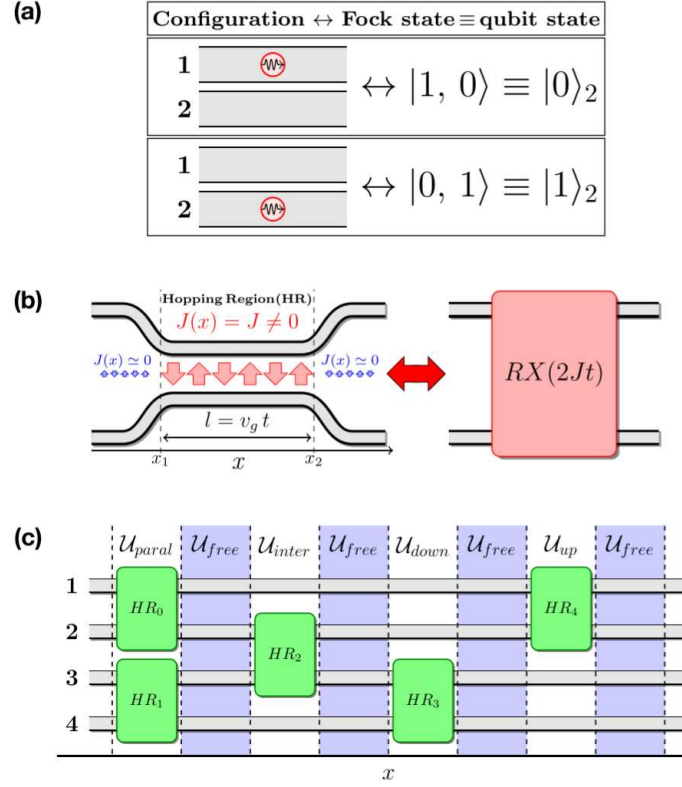


Figure 3.9: Photonic QC elements. **a)** Qubit logical state obtained with a dual-rail encoding of a pair of photonic channels, in which single-photon Fock states can be injected into each waveguide and measured at the output through single-photon detectors. **b)** The hopping process of a single photon between two channels can be viewed as an RX gate (see paragraph Single-qubit gates). **c)** Single block of a two-qubit information processing unit: Multiple repetitions of such an elementary block can be concatenated to obtain the desired two-qubit gate. Hopping regions (HR) (green boxes) depend on two parameters (J_m, t_m) , while free propagation (FP) (blue regions) depends only on a single parameter (t_m) for a given value of the photon-photon nonlinearity, U . The FP layers within the U_{inter} , U_{down} , and U_{up} regions are assumed to have fixed propagation time set by the corresponding HR terms (implicitly represented by a simple line in the sketch). Picture taken from the original in Ref. [17].

with $\mathbb{I}$ and σ_x being the identity and the Pauli-X matrices respectively. In the *free propagation regions* (FP), the channels are decoupled ($J = 0$) and each photon propagates independently, accumulating a nonlinear phase shift Γ if two photons are simultaneously present in the same channel. Indeed, since for a single propagating channel, the Hamiltonian H is diagonal, one has the following evolution matrix

$$U_{FP} = e^{-iHt} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & e^{-i\omega t} & 0 \\ 0 & 0 & e^{-i2(\Gamma+\omega)t} \end{bmatrix}. \quad (3.9)$$

In particular the unitaries are:

- U_{free} : all four channels in free propagation, no hopping; the self-Kerr nonlinearity Γ acts on any two-photon state present in the same channel.

$$U_{free} = U_{FP} \otimes U_{FP} \otimes U_{FP} \otimes U_{FP}$$

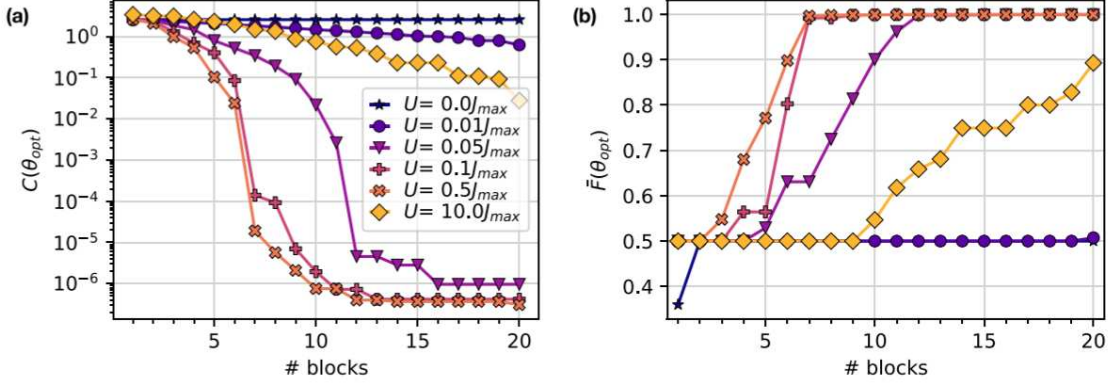


Figure 3.10: CNOT optimization. Comparison of **a)** the best cost function (see Eq.3.12) and **b)** average gate fidelity (see Eq.3.13) for different values of the Kerr nonlinearity U (in hopping parameter J_{max} units) after the optimization of the block structure. The free propagation and interaction times were both fixed to $t = 1$ in dimensionless units. Picture from [17].

- $\mathcal{U}_{\text{paral}}$: simultaneous hopping between channels (1,2) and (3,4); acts independently on the two qubits.

$$\mathcal{U}_{\text{paral}} = U_{\text{HR}} \otimes U_{\text{HR}}$$

- $\mathcal{U}_{\text{inter}}$: hopping between channels (2,3); this is the only operation that directly couples the two qubits and, in the presence of Γ , is responsible for generating entanglement.

$$\mathcal{U}_{\text{inter}} = U_{\text{FP}} \otimes U_{\text{HR}} \otimes U_{\text{FP}}$$

- $\mathcal{U}_{\text{up}}, \mathcal{U}_{\text{down}}$: hopping on channels (1,2) or (3,4) respectively, with the remaining channels in free propagation.

$$\mathcal{U}_{\text{down}} = U_{\text{FP}} \otimes U_{\text{FP}} \otimes U_{\text{HR}}$$

$$\mathcal{U}_{\text{up}} = U_{\text{HR}} \otimes U_{\text{FP}} \otimes U_{\text{FP}}$$

The hopping parameters $\{J_{ij}\}$ in each block are the free variables of the optimization (denoted as $\{\theta_s\}$) while Γ and the propagation times are fixed. The total unitary is

$$U_{\text{tot}}(\{\theta_s\}) = \prod_{b=1}^M U_b = U_M U_{M-1} \cdots U_1. \quad (3.10)$$

By increasing the number of blocks M , the circuit becomes more expressive and the optimizer can find configurations that more faithfully reproduce the target gate.

CNOT gate optimization.

The optimization targets the CNOT gate, whose ideal matrix representation on the computational basis $\mathcal{S}$ is (see eq:1.6)

$$T_{\text{CNOT}} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix}. \quad (3.11)$$

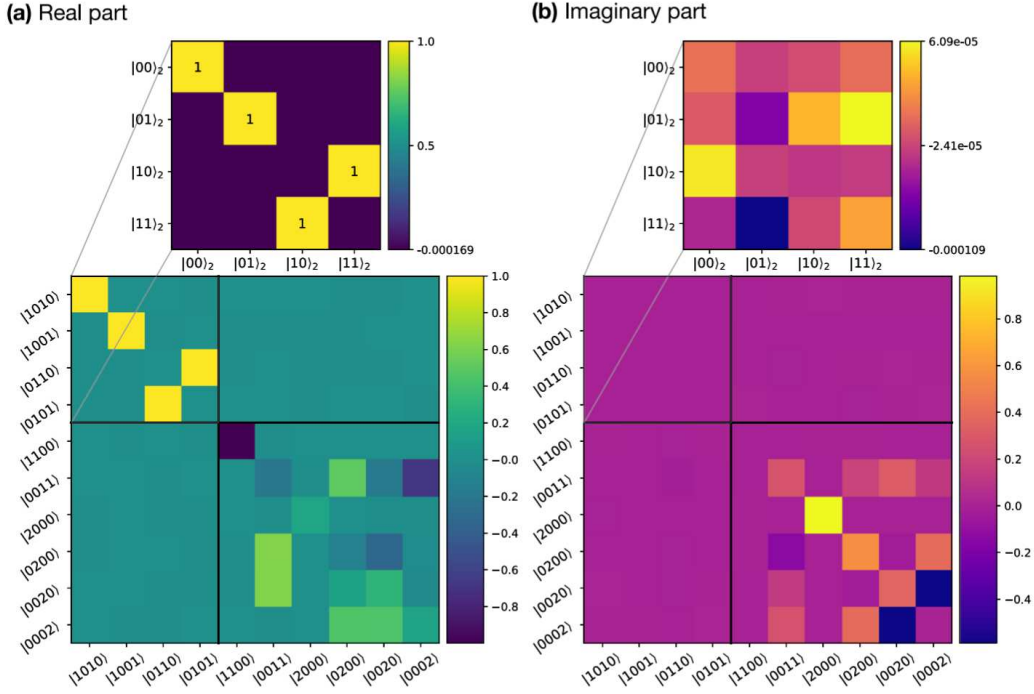


Figure 3.11: Best CNOT matrix. **a)** Real and **b)** imaginary parts (rounded at the third decimal digit) of the approximate CNOT matrix on the basis of all the possible twophoton input/output states obtained by using a 12-block structure with $U = 0.5J_{max}$ and $\bar{F}(\{\theta_s^{\text{opt}}\}) \approx 99.96\%$. Vertical and horizontal black lines divide the logic space from the one based on states that are not on the computational basis. The projection of logic states of the computational space is negligible, as clearly seen from the plot. The zoom highlights the gate matrix in the two-qubit logic space. Picture from [17].

The *cost function* minimized at each step is the Frobenius norm between the optimized unitary $\tilde{U}_{\text{tot}}$ and the target T_{CNOT} ,

$$C(\{\theta_s\}) = \|\tilde{U}_{\text{tot}}(\{\theta_s\}) - T_{\text{CNOT}}\|_F^2, \quad (3.12)$$

giving a set of optimized parameters θ_s^{opt} , and the performance is quantified by the *average gate fidelity*

$$\bar{F}(\{\theta_s^{\text{opt}}\}) = \frac{1}{|\mathcal{S}|} \sum_{i \in \mathcal{S}} |\langle i | \tilde{U}_{\text{tot}}^\dagger(\{\theta_s^{\text{opt}}\}) T | i \rangle|^2. \quad (3.13)$$

The results are shown in Fig. 3.10: for nonlinearities in the range $0.05 \leq \Gamma/J_{\text{max}} \leq 0.5$, the fidelity converges rapidly towards unity as the number of blocks increases, reaching $\bar{F} \approx 99.96\%$ with just 12 blocks at $\Gamma/J_{\text{max}} = 0.5$. For negligible nonlinearity the cost function cannot be minimized, consistently with the known impossibility of deterministic two-qubit gates in linear optics.

The real and imaginary parts of the optimized gate matrix are shown in Fig. 3.11, confirming that no logical state leaks outside the computational basis, a direct signature of the deterministic nature of the gate.

Conclusion.

The results obtained by Scala et al. [17] demonstrate that deterministic two-qubit entangling gates are achievable in nonlinear photonic interferometers with dual-rail encoding, overcoming the fundamental probabilistic limitation of linear optics for quantum computing applications. The optimal CNOT gate reaches a theoretical fidelity of 99.96% with 12 elementary blocks at $\Gamma/J_{\max} = 0.5$, a value that compares favourably with the state-of-the-art CNOT fidelities in superconducting circuits (99.77%) and trapped-ion devices (99.4%). The gate is deterministic in the sense that no logical state leaks outside the computational basis at the output, and its structure is preserved even in the presence of realistic losses and fabrication imperfections, as verified through open quantum system simulations.

On the other hand, the practical scalability of the scheme remains an open challenge. The required concatenation of 10–20 elementary blocks implies a total circuit depth that, combined with the inevitable photon losses in realistic waveguides, may limit the achievable fidelity in actual devices. Furthermore, the scheme requires on-chip self-Kerr nonlinearities at the single-photon level, which, while compatible with polariton platforms and potentially with silicon nitride (SiN) or lithium niobate-on-insulator (LNOI) technologies, have not yet been demonstrated with sufficient strength and control in a propagating geometry.

Summary

In this Chapter, we have summarized the recent developments in the quest for the implementation of two-qubit quantum operations in photonics: both Lahini et al. [15] and Scala et al. [17] have demonstrated, essentially, that a deterministic CNOT gate is achievable beyond the linear optics paradigm, although in conceptually different ways and with different trade-offs.

In the quantum walk approach by Lahini et al., determinism is achieved by exploiting the on-site interaction Γ between ultra-cold bosonic atoms in an optical lattice. The proposed gate implementation requires only four lattice sites and two atoms, with a static Hamiltonian optimized once, a remarkably compact and conceptually clean architecture. In the quantum walk approach, the interaction is distributed continuously along the propagation direction, rather than localized in discrete hopping regions as in Scala et al. The fidelity of 99.6% is limited only by the experimental bounds imposed on the tunneling and interaction parameters, and approaches unity as these bounds are relaxed. However, the platform is inherently atomic: it requires single-site addressing, single-atom detection, and precise tuning of the on-site interaction strength Γ , which in atomic systems is controlled via external magnetic fields through Feshbach resonances — all of which are experimentally demanding and difficult to scale to many qubits.

On the other hand, in the nonlinear interferometer approach of Scala et al., determinism is instead achieved in a photonic platform by exploiting a distributed self-Kerr nonlinearity in a concatenated multi-block circuit. The fidelity of 99.96% is higher, and the scheme is general (it does not rely on a specific material and is in principle compatible with several photonic platforms). On the other hand, the circuit requires 10–20 elementary blocks, implying a significantly larger physical footprint and a higher sensitivity to propagation losses compared to the four-site lattice of Lahini.

In summary, the quantum walk architecture offers compactness and conceptual simplicity at the cost of experimental complexity in the atomic platform, while the nonlinear interferometer achieves higher fidelity and platform flexibility at the cost of a more complex circuit structure. The natural question is whether these two strengths can be combined. That is, whether a nonlinear quantum walk in a photonic platform could yield a deterministic CNOT gate that is both compact and experimentally accessible. This is precisely the question addressed in the original work presented in the following Chapter.

Chapter 4

Algorithmic optimization of photonic two-qubit gates

The results presented in Chapters 2 and 3 establish a clear picture highlighting where photonic quantum computing currently stands, and what its main bottlenecks are. On one side, the linear-optical paradigm, culminating in the quantum-walk-based optimized CNOT gate implementation from Lahini et al. [15], and its experimental realization by Chapman et al. [16], achieves the fundamental bound of success probability $p_{\text{succ}} = 1/9$ with a remarkably compact architecture: six straight, parallel waveguides, no ancilla photons, straight-line routing. More than that, it confirms the previous theoretical bound by Ralph et al., proposed in a minimal interferometric architecture [10]. On the other side, the nonlinear interferometer scheme proposed by Scala et al. [17] demonstrates that deterministic two-qubit gates are in principle achievable in photonic platforms, provided that the circuit is complemented by a distributed self-Kerr nonlinearity, although at the cost of a multi-block concatenated architecture with a significantly larger physical footprint.

At this point of our tale, these two results naturally bring to ask ourselves a question: *can the architectural simplicity of the quantum-walk approach be combined with the power of a nonlinear photon-photon interaction to push the success probability beyond the $1/9$ bound, or even to reach deterministic operation, while keeping the minimal six-waveguide geometry?*

In fact, this chapter addresses this question directly. First, we take the six-site waveguide lattice of Lahini et al. as our starting point: same Hamiltonian structure, dual-rail encoding, fixed propagation time ($T = 1$). To go beyond this setting, in the final part of the work we introduce a uniform on-site self-Kerr nonlinearity (parametrized by Γ) acting on all sites simultaneously and throughout propagation. Unlike the scheme of Scala et al., where the circuit is structured as a concatenation of multiple blocks with piecewise-constant hopping parameters, here the Hamiltonian terms are always on: both the hopping between waveguides and the self-Kerr phase accumulation act continuously and simultaneously over the full propagation length. Furthermore, while Scala et al. assumed identical waveguides with uniform on-site energies, here each waveguide is allowed to have a distinct on-site energy E_m , which is treated as a free parameter in the optimization, along the same spirit already used by the Lahini et al [15].

This combination, the architectural simplicity of the parallel waveguide geometry, enriched with individual on-site energies and the continuous distributed nonlinearity, is

what allows us to find, in the same minimal six-waveguide structure, a regime that goes beyond the previously found probabilistic bound, and approaches the deterministic operation previously demonstrated but now in a more compact architecture. Practical realization of this paradigm will be discussed in the end: in fact, it is shown how this regime could be naturally realized by photons propagating in a waveguide array with an intrinsic material nonlinearity, such as exciton-polariton systems or silicon nitride waveguides.

Throughout this work, we systematically explore the trade-off between gate fidelity and success probability as a function of Γ by means of a numerical optimization of the ten free parameters of the Hamiltonian: five on-site energies and five nearest-neighbour hopping coefficients, over a dense grid of target success probabilities. In the linear case ($\Gamma = 0$), our optimization results essentially reproduce the ones previously obtained from Lahini et al., thus confirming that the fundamental $1/9$ bound is saturated, and that no configuration in this geometry can exceed it. On the other hand, we numerically find that as Γ is increased, new regions of the fidelity-probability trade-off become accessible: high fidelity is achieved at success probabilities well above $1/9$, a regime that is strictly forbidden by the no-go results of linear optics. For sufficiently large nonlinearity the optimization approaches fully deterministic operation, with fidelity close to unity at $p_{\text{succ}} \rightarrow 1$. The latter are the main original results of this numerical campaign, and of the thesis.

4.1 The nonlinear quantum walk

4.1.1 Hamiltonian and waveguide lattice

The physical system studied in this work is a one-dimensional array of six evanescently coupled straight waveguides (see, e.g., Fig. 4.1). Eventually, we will allow for self-Kerr nonlinearity to be present in each waveguide during propagation, which is compatible with the definition of a *nonlinear quantum walk* in a one-dimensional waveguide array. As it is typical of similar schemes in LOQC, the outer waveguides, $|a_1\rangle$ and $|a_2\rangle$, are empty ancillary modes containing no photons, while the four inner waveguides encode the control and target qubits in pairs: $|c_1\rangle, |c_2\rangle$ represent the control qubit and $|t_1\rangle, |t_2\rangle$ represent the target qubit, following dual-rail encoding. This model is then described by the Hamiltonian

$$\mathcal{H} = \sum_{m=1}^6 E_m a_m^\dagger a_m + \sum_{m=1}^5 J_{m,m+1} (a_m^\dagger a_{m+1} + a_{m+1}^\dagger a_m) + \Gamma \sum_{m=1}^6 a_m^\dagger a_m (a_m^\dagger a_m - 1), \quad (4.1)$$

in which $a_m^\dagger$ and a_m are the bosonic creation and annihilation operators for a photon in the m -th waveguide, satisfying $[a_m, a_{m'}^\dagger] = \delta_{mm'}$. As already mentioned in the previous Chapter, this Hamiltonian can also be viewed as the standard Bose-Hubbard model, in which each term has a direct and well justified physical origin.

The first term assigns an *on-site energy* E_m to each photon present in waveguide m , weighted by the number operator $\hat{n}_m = a_m^\dagger a_m$; physically, E_m is proportional to the effective refractive index of waveguide m relative to a reference value, and can be tuned by adjusting the waveguide width or the local refractive index contrast.

The second term describes *nearest-neighbor hopping*: $a_{m+1}^\dagger a_m$ annihilates a photon in

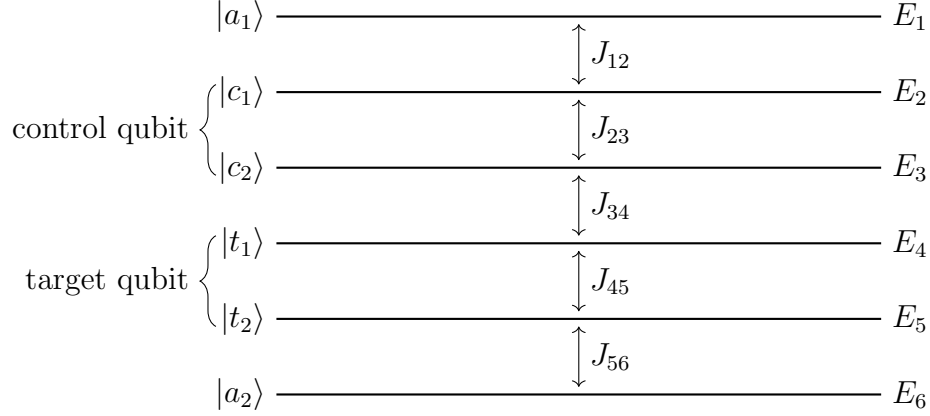


Figure 4.1: The six waveguide array employed to implement a nonlinear quantum walk with two-qubit encoding and ancillary channels.

waveguide m and creates one in waveguide $m + 1$, modelling the evanescent tunnel coupling between adjacent waveguide channels. The coefficients $J_{m,m+1}$ are then proportional to the evanescent coupling strength, which physically depends exponentially on the waveguide separation: a smaller gap yields a larger $|J_{m,m+1}|$. The term is restricted to nearest neighbours since the coupling decays exponentially with distance, and it is written together with its Hermitian conjugate ($a_m^\dagger a_{m+1}$) to ensure the Hamiltonian is Hermitian. Positive and negative values of $J_{m,m+1}$ correspond to different relative phases of the coupling, and both are in principle accessible by engineering the waveguide geometry. In the algorithmic optimization presented in the following, both E_m and $J_{m,m+1}$ will be allowed to take any real value within a symmetric interval $[-E_{\max}, E_{\max}]$.

The third term is the *on-site self-Kerr* interaction introduced before: the product $\hat{n}_m(\hat{n}_m - 1)$ vanishes whenever $\hat{n}_m = 0$ or 1 , and equals 2 only when two photons simultaneously occupy the same site, correctly capturing the fact that the nonlinearity is a pairwise interaction effect. This is the standard second-quantized form for a $\chi^{(3)}$ nonlinear interaction energy, with Γ proportional to the third-order nonlinear susceptibility of the waveguide material.

In the photonic realization of this scheme, the longitudinal coordinate along the waveguides, e.g., z , plays the role of time, and the system evolves according to the unitary operator

$$U = e^{-i\pi\mathcal{H}T}, \quad (4.2)$$

in which Hamiltonian parameters are given in dimensionless units (the evolution is regulated by phases that are units of π , see the exponent in Eq. 4.2), and T is the dimensionless propagation time. In the following, we will fix $T = 1$ throughout the optimization procedure. In fact, the choice for a fixed $T = 1$ is not arbitrary, but it is motivated by a consistency check against the physical scales of the photonic device. Let us briefly discuss this point here, once for all.

Let us assume, for concreteness, to have an array made of waveguides of length $L = 60 \mu\text{m}$, and that photonic signals propagate in the single waveguide at a group velocity $v_g = 30 \mu\text{m}/\text{ps}$ (estimated, e.g., from unpublished experimental data, and reported in a previous work on exciton-polariton waveguides [69]). Then, the physical propagation

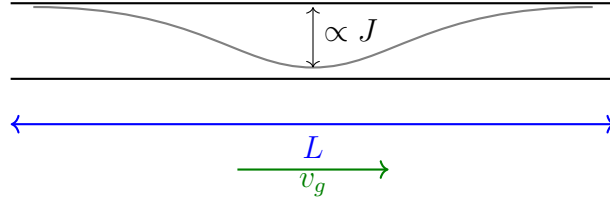


Figure 4.2: Schematic illustration of a single amplitude oscillation of the photon field propagating in a waveguide of length L at group velocity v_g .

time would be $T_{\text{phys}} = L/v_g \simeq 2$ ps. Requiring the coupling strength J be large enough to produce a full oscillation of the excitation over this propagation length, i.e., $JT_{\text{phys}} \sim \pi$, yields a coupling energy $\hbar J \simeq 1$ meV, in agreement with the order of magnitude reported experimentally. This estimate confirms that couplings on the order of π are physically reasonable, and it is this scale that ultimately justifies fixing the dimensionless propagation time to $T = 1$ in the model. This choice sets the natural unit of the problem: all energies and hopping rates are expressed in units of $1/\pi$.

Having fixed the propagation time this way, we have then verified numerically that the model parameters obtained from the optimization scale consistently with Eq. (4.2): while the optimized parameters are of order 10^0 for $T = 1$, they reduce to order 10^{-1} for $T = 10$. This behaviour confirms that the model correctly captures the trade-off between propagation time and interaction strength, i.e., $JT_{\text{phys}} \sim \pi$. This consideration ultimately supports the choice of working at $T = 1$ throughout, since it keeps the optimized parameters within a numerically convenient range.

This physical estimate also motivates the choice for the bounds of the parameters to be optimized. In the linear case ($\Gamma = 0$), an energy scale of order π is expected to be sufficient to induce a full oscillation in the circuit, consistent with the experimental estimate above; therefore, the bounds were chosen as $[-4, 4]$, which comfortably contains this scale. When $\Gamma \neq 0$, the nonlinear term increases the energy scale required to achieve the same dynamics, so the bounds must be widened accordingly as Γ increases, in order to avoid artificially constraining the optimizer.

The on-site energy of the first waveguide is fixed to $E_1 = 0$ without loss of generality, leaving ten free parameters: $E_2, \dots, E_6$ and $J_{12}, \dots, J_{56}$.

Finally, it is important to contrast the role of Γ in this system with its role in the two reference works. In Lahini et al., the nonlinear case ($\Gamma \neq 0$) was studied for ultra-cold bosonic atoms in optical lattices, where the on-site interaction is strong and localized, and the photonic case was treated as strictly linear ($\Gamma = 0$). Unlike the scheme of Scala et al., where the circuit is structured as a concatenation of multiple blocks with piecewise-constant hopping parameters and the nonlinearity accumulates over many sequential segments, here, the system consists of a single propagation region of fixed length in which hopping and self-Kerr nonlinearity act simultaneously and continuously, with all parameters optimized globally.

4.1.2 Qubit encoding

Logical qubits are encoded, as usual, by applying the dual-rail scheme, in which a single photon distributed across two adjacent waveguides represents one qubit. Formally, the

logical states of a single qubit are thus expressed as

$$|0\rangle_{dr} = a_1^\dagger |\Omega\rangle = |1\rangle \otimes |0\rangle = |1, 0\rangle \quad |1\rangle_{dr} = a_2^\dagger |\Omega\rangle = |0\rangle \otimes |1\rangle = |0, 1\rangle, \quad (4.3)$$

in which $|\Omega\rangle = |0, 0\rangle$ is the vacuum state and the labels 1, 2 refer to the left and right waveguide of the pair. For a two-qubit system, two photons are injected into the lattice, one per qubit. The four computational basis states are encoded in waveguides 2–5, with waveguides 1 and 6 serving as auxiliary vacuum modes:

$$\begin{aligned} |00\rangle &= a_2^\dagger a_4^\dagger |\Omega\rangle, \\ |01\rangle &= a_2^\dagger a_5^\dagger |\Omega\rangle, \\ |10\rangle &= a_3^\dagger a_4^\dagger |\Omega\rangle, \\ |11\rangle &= a_3^\dagger a_5^\dagger |\Omega\rangle. \end{aligned} \quad (4.4)$$

The first qubit (control) is encoded in waveguides 2–3, the second qubit (target) in waveguides 4–5. Waveguides 1 and 6 are implemented empty and act as vacuum ancilla modes, as in the original scheme from Ralph et al. [10]. The full Fock space of the system is spanned by states with minimum zero photon and at most two photons distributed across six modes; to correctly capture all two-photon states, each mode is truncated to a maximum occupation of 2 photons, yielding a local Hilbert space of dimension 3 (i.e., accounting for photon occupations 0, 1, and 2). The total Hilbert space therefore has dimension $3^6 = 729$.

The action of the device on the logical subspace is described by a 4×4 matrix $G_{4 \times 4}$, obtained by projecting the full unitary U onto the logical basis:

$$(G_{4 \times 4})_{ij} = \langle \phi_i | U | \phi_j \rangle, \quad (4.5)$$

in which $\{|\phi_j\rangle\}$ are the four logical basis states defined in Eq. (4.4).

4.2 Definitions of Success probability and Fidelity

4.2.1 Success probability

In the linear case ($\Gamma = 0$), the matrix $G_{4 \times 4}$ is not unitary on the logical subspace: amplitude leaks into states outside the computational basis (those with both photons in the same waveguide pair, or in the ancilla modes), and the gate is necessarily probabilistic. To justify the definition of success probability used in this work, it is useful to recall the general formalism of quantum measurements, following Nielsen and Chuang [1].

First, we remind that a general quantum measurement is described by a collection of measurement operators, $\{M_m\}$, acting on the state space of the system, where the index m labels the possible outcomes. If the system is in state $|\psi\rangle$ immediately before the measurement, the probability that outcome m occurs is

$$p(m) = \langle \psi | M_m^\dagger M_m | \psi \rangle, \quad (4.6)$$

and the measurement operators satisfy the completeness relation (see 1.5):

$$\sum_m M_m^\dagger M_m = \mathbb{I}. \quad (4.7)$$

This relation guarantees that probabilities sum up to one over all possible outcomes, $\sum_m p(m) = 1$, and this is the only constraint imposed on the operators M_m : unlike projective measurements, the M_m need not be orthogonal projectors, which makes this formalism general enough to describe measurements with intrinsic loss, such as postselection on a subspace.

In our setting, the relevant outcome is binary: either the two output photons are detected in the logical subspace spanned by $\{|\phi_j\rangle\}$ (success), or they are found elsewhere in the Hilbert space, e.g. bunched in the same waveguide or leaked into the ancilla modes (failure). The operator describing a successful event is the projection of the full unitary evolution U onto the logical subspace, which is precisely $G_{4\times 4}$ as defined in Eq. (4.5). Identifying $M_{\text{succ}} \equiv G_{4\times 4}$ in Eq. (4.6), the probability of success for a given logical input state $|\phi_j\rangle$ is

$$p_{\text{succ}}^{(j)} = \langle \phi_j | G_{4\times 4}^\dagger G_{4\times 4} | \phi_j \rangle = \left(G_{4\times 4}^\dagger G_{4\times 4} \right)_{jj}, \quad (4.8)$$

i.e. the j -th diagonal entry of $G_{4\times 4}^\dagger G_{4\times 4}$ in the logical basis. Since the optimization is required to perform well on all four logical inputs simultaneously, and not only on a particular one, the natural figure of merit is the average of $p_{\text{succ}}^{(j)}$ over the four basis states, weighted uniformly. This average is exactly the normalized trace of $G_{4\times 4}^\dagger G_{4\times 4}$, which is equivalent to the mean of its eigenvalues (equivalently, the mean of the squared singular values σ_k^2 of $G_{4\times 4}$, since $G_{4\times 4}^\dagger G_{4\times 4}$ is Hermitian and positive semi-definite):

$$p_{\text{succ}} = \frac{1}{4} \sum_{j=1}^4 p_{\text{succ}}^{(j)} = \frac{1}{4} \text{Tr} \left(G_{4\times 4}^\dagger G_{4\times 4} \right) = \frac{1}{4} \sum_{k=1}^4 \sigma_k^2. \quad (4.9)$$

This definition has a clear physical interpretation. Each singular value $\sigma_k^2 \in [0, 1]$ measures the fraction of the k -th input amplitude that survives projection onto the logical output subspace; averaging over k gives the overall fraction of two-photon events that are correctly detected in the computational basis, regardless of which of the four logical states was injected.

In the deterministic gating limit (i.e., $p_{\text{succ}} \rightarrow 1$), all the output amplitudes remain within the logical subspace: $G_{4\times 4}$ becomes proportional to a unitary matrix, so that every $\sigma_k = 1$ and $p_{\text{succ}} = \frac{1}{4} \text{Tr}(G_{4\times 4}^\dagger G_{4\times 4}) = 1$, meaning no amplitude leaks outside the computational basis.

Conversely, the condition $p_{\text{succ}} = 0$ would describe a gate whose output never overlaps with any logical state of the computational basis of the two qubits. Practically, this never occurs in practice, since some residual probability always remains in the computational basis even for unoptimized parameters.

4.2.2 Two-qubit gate fidelity

While p_{succ} quantifies how often the two-qubit gate succeeds, it carries no information about whether the output, conditioned on success, actually implements the desired CNOT transformation. This second requirement is captured by the *gate fidelity*.

In general terms, the standard notion of fidelity between two pure quantum states, $|\psi\rangle$ and $|\varphi\rangle$, is given by the overlap between the two, as introduced by Nielsen and Chuang [1]

$$F(|\psi\rangle, |\varphi\rangle) = |\langle \psi | \varphi \rangle|, \quad (4.10)$$

which equals 1 if and only if the two states coincide up to a global phase, and 0 if they are orthogonal. When the objects to compare are not states but unitary operators, as it is our case with the ultimate goal of comparing the optimized gate (i.e., the matrix $G_{4 \times 4}$) with the target operation (i.e., the two-qubit matrix U_{CNOT}), this notion must be lifted from the level of vectors to the level of operators. In this context, Nielsen and Chuang outline this generalization in their discussion of the diamond norm and related distance measures between quantum operations, where the central idea is to treat an operator as a vector in the space of linear operators on the Hilbert space, equipped with the Hilbert-Schmidt inner product

$$\langle A, B \rangle_{\text{HS}} \equiv \text{Tr}(A^\dagger B). \quad (4.11)$$

This inner product turns the space of $d \times d$ operators into a d^2 -dimensional Hilbert space, in which the natural notion of overlap between two operators A and B is precisely the analogue of Eq. (4.10).

A fully explicit construction along these lines is given by Wang, Yu, and Yi, who define an *operator fidelity* between two operators A and B by first normalizing them with respect to the Hilbert-Schmidt norm, $A/\sqrt{\text{Tr}(AA^\dagger)}$ and $B/\sqrt{\text{Tr}(BB^\dagger)}$, and then taking the normalized overlap [72]

$$F(A, B) = \frac{|\text{Tr}(A^\dagger B)|}{\sqrt{\text{Tr}(AA^\dagger) \text{Tr}(BB^\dagger)}}. \quad (4.12)$$

This definition satisfies the key properties expected from a fidelity measure (i.e., normalization, symmetry, and invariance under unitary transformations), reducing, when A and B are pure-state projectors, to the usual state fidelity defined in Eq. (4.10). When $A = U_0$ and $B = U_1$ are both d -dimensional unitary operators, $\text{Tr}(U_0 U_0^\dagger) = \text{Tr}(U_1 U_1^\dagger) = d$, and Eq. (4.12) reduces to the particularly simple form

$$F(U_0, U_1) = \frac{1}{d} \left| \text{Tr}(U_0^\dagger U_1) \right|. \quad (4.13)$$

In our case, the comparison is not between two unitary operators of the same dimension, but rather between the ideal $d = 4$ unitary U_{CNOT} and the matrix $G_{4 \times 4}$, which is generally *not* unitary, since amplitude is lost outside the logical subspace whenever $p_{\text{succ}} < 1$. To apply Eq. (4.12) consistently, $G_{4 \times 4}$ must first be normalized with respect to its own Hilbert-Schmidt norm, exactly as prescribed by Wang et al. for a generic (not necessarily unitary) operator:

$$\tilde{G} = \frac{G_{4 \times 4}}{\|G_{4 \times 4}\|_F}, \quad \|G_{4 \times 4}\|_F \equiv \sqrt{\text{Tr}(G_{4 \times 4}^\dagger G_{4 \times 4})}, \quad (4.14)$$

where $\|\cdot\|_F$ denotes the Frobenius norm, which coincides with the Hilbert-Schmidt norm $\sqrt{\text{Tr}(AA^\dagger)}$ used in Eq. (4.12). Since U_{CNOT} is already unitary,

$$\|U_{\text{CNOT}}\|_F = \sqrt{\text{Tr}(U_{\text{CNOT}}^\dagger U_{\text{CNOT}})} = \sqrt{4} = 2$$

is a fixed normalization constant. Substituting $A = U_{\text{CNOT}}$ and $B = \tilde{G}$ into Eq. (4.12), and using $\text{Tr}(U_{\text{CNOT}} U_{\text{CNOT}}^\dagger) = 4$ together with $\text{Tr}(\tilde{G} \tilde{G}^\dagger) = 1$ by construction, gives the

fidelity used throughout this work:

$$\mathcal{F} = \frac{\left| \text{Tr} \left(U_{\text{CNOT}}^\dagger \tilde{G} \right) \right|}{\|U_{\text{CNOT}}\|_F}, \quad (4.15)$$

where the absolute value makes $\mathcal{F}$ invariant under multiplication of $\tilde{G}$ by an arbitrary global phase, consistent with the properties of the operator fidelity discussed in [72] and with the phase-invariance of the state fidelity in Eq. (4.10).

By construction, $\mathcal{F} \in [0, 1]$, with $\mathcal{F} = 1$ corresponding to a perfect CNOT gate up to a global phase, and $\mathcal{F} = 0$ corresponding to an output completely orthogonal to the target operation in the Hilbert-Schmidt sense. This definition can be interpreted as a lower bound on the average gate fidelity over all input states in the logical subspace, since it only tests the overlap of $G_{4 \times 4}$ with U_{CNOT} as a whole, rather than averaging fidelities computed independently on each of the four logical basis states.

We remind the the target CNOT matrix in the computational basis $\{|00\rangle, |01\rangle, |10\rangle, |11\rangle\}$ is

$$U_{\text{CNOT}} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix}, \quad (4.16)$$

with $\|U_{\text{CNOT}}\|_F = 2$.

Together, $\mathcal{F}$ and p_{succ} characterize the performance of the two-qubit gate: a useful device must achieve high fidelity (the output state is close to the target) and simultaneously a success probability that is as large as possible (the gate succeeds frequently). These two quantities define a Pareto front in the $(\mathcal{F}, p_{\text{succ}})$ plane, that is, the set of multi-objective optimal solutions for which no feasible solution can improve one objective (such as, e.g., the Fidelity) without degrading the other (such as, e.g., the success probability). The goal of the optimization is to trace this front and analyse its evolution as a function of fixed values of the nonlinearity, Γ .

4.3 Optimization strategy

The figures of merit introduced in the previous section, fidelity and success probability, are in general competing quantities: increasing one typically comes at the cost of the other. The goal of the optimization is therefore not to find a single optimal configuration of the physical parameters, but to reconstruct the trade-off between the two quantities, i.e. the *Pareto front* of the problem, as a function of the nonlinearity Γ .

4.3.1 Parameters, search space and cost function

The free parameters of the model are the on-site energies E_i and the tunneling rates $J_{i,i+1}$ of the six-site lattice, for a total of $n = 10$ independent parameters

$$\boldsymbol{\theta} = (E_2, \dots, E_6, J_{12}, \dots, J_{56}),$$

with $E_1 = 0$ fixed as a reference, as discussed in the previous section. Each parameter is bounded in a symmetric interval $[-E_{\text{max}}, E_{\text{max}}]$, wide enough to allow full exploration

of the relevant interference regimes while keeping the search space bounded, as required by the optimization algorithm.

For a given choice of parameters $\boldsymbol{\theta}$, the gate $G_{4 \times 4}$ is obtained by evolving the logical basis under the Hamiltonian $H(\boldsymbol{\theta})$ for the fixed time $T = 1$, as described above, and the fidelity $\mathcal{F}(\boldsymbol{\theta})$ and success probability $p_{\text{succ}}(\boldsymbol{\theta})$ follow from Eqs. (4.9) and (4.15). To reconstruct the Pareto front between these two quantities, we adopt the so called *weighted-constraint scalarization method*, it is a multi-objective optimization technique that transforms multiple conflicting objectives into a single-objective problem. So, we fix a target success probability p_{target} and minimize the infidelity while penalizing deviations from the target, via the cost function

$$C(\boldsymbol{\theta}) = (1 - \mathcal{F}(\boldsymbol{\theta})^2) + w (p_{\text{succ}}(\boldsymbol{\theta}) - p_{\text{target}})^2, \quad (4.17)$$

where w is a penalty weight chosen to strongly enforce the constraint on p_{succ} relative to the infidelity term. Repeating the minimization of $C(\boldsymbol{\theta})$ for a set of 80 target values p_{target} linearly spaced in $[0.02, 0.99]$ yields, for each target, the best achievable fidelity compatible with that constraint; the resulting set of points $(\mathcal{F}^*, p_{\text{succ}}^*)$ traces out the Pareto front of the problem for the chosen value of Γ .

4.3.2 Optimization algorithm and strategy

The cost function defined in Eq. (4.17) depends on the parameters through the matrix exponential $e^{-i\pi T H(\boldsymbol{\theta})}$, whose analytical gradient with respect to $\boldsymbol{\theta}$ is not implemented; moreover, the underlying physical problem is expected to exhibit multiple local optima, owing to the interferometric, oscillatory dependence of the gate on the tunneling and on-site energy parameters. For these reasons, a *derivative-free, global* optimization strategy is required, as opposed to gradient-based local methods (e.g. BFGS), which would normally need an explicit analytical gradient and would be predisposed to converging to suboptimal local minima close to the chosen starting point. Several global, derivative-free optimization algorithms were tested before settling on the final scheme: Differential Evolution and Dual Annealing from the `scipy.optimize` package, and ISRES and AUGLAG combined with COBYLA from the `NLopt` library¹, together with L-BFGS-B as a gradient-based local baseline, in which the gradient was approximated via finite differences in the absence of an analytical expression; this approach, however, remains computationally expensive in higher dimensions and, being a local method, is still prone to convergence toward suboptimal minima. All of these methods converged, qualitatively, to the same overall trade-off between fidelity and success probability, which is itself a useful consistency check: it indicates that the resulting Pareto front reflects a genuine feature of the physical system, rather than an artifact of a specific optimizer. However, at comparable computational budget, all of these alternatives produced visibly noisier fronts, with larger point-to-point fluctuations in the optimal fidelity between neighbouring target probabilities. The scheme finally adopted, described below, consistently produced the smoothest and most consistent front among those tested, and was therefore used for all the results presented in this chapter.

The adopted optimization strategy is a two-stage scheme, implemented via the `NLopt` library. In the first stage, a global search is carried out using the *Multi-Level*

¹`NLopt` source website: <https://nlopt.readthedocs.io/en/latest/>

Single-Linkage algorithm (MLSL) with a low-discrepancy sampling sequence (LDS). MLSL is a stochastic global optimization meta-algorithm: it samples candidate starting points over the parameter space and launches a local optimizer from a selected subset of them, discarding points that fall too close to previously identified local minima (the so-called *clustering* criterion), so as to avoid redundant local searches. This guarantees, with probability tending to one as the number of function evaluations increases, convergence to the global minimum. The local optimizer used internally by MLSL is BOBYQA (Bound Optimization BY Quadratic Approximation), a derivative-free trust-region method that builds a quadratic interpolation model of the objective function from sampled values and uses it to determine the search direction, without requiring gradient information. During the global stage, BOBYQA is run with a loose convergence tolerance of 10^{-3} and a small evaluation budget per local search, 20 iterations, so as to explore many regions of the parameter space efficiently rather than converging to high precision in any single one, we will refine the local optimization in the next stage. The number of iterations performed by the global optimizer for each point on the Pareto front ranged from 11'000, in the linear case, to 18'000, for $\Gamma = 1$, with convergence tolerance of 10^{-3} . This increase was necessary because a larger nonlinearity widens the search bound, requiring more iterations to properly explore the parameter space. In the second stage, the best point found by the global search is refined with a separate high-precision run of BOBYQA, with a tight convergence tolerance of 10^{-6} , and a larger evaluation budget compared to the local refinement of the previous stage, between 1'000 and 1'500 iterations, to obtain the final optimized parameters. This combination of a multi-start global search (MLSL) with a local, gradient-free refinement step (BOBYQA) is a common choice for optimizing variational photonic systems with nonlinear components, having also been used, for example, to train quantum optical neural networks [18].

For each target probability, p_{target} , the full two-stage procedure described above is repeated from scratch, i.e., a new MLSL global search over the entire parameter space is performed, rather than warm-starting the optimization from the optimal parameters found for a neighbouring target value. This choice is deliberate: reusing the previous optimum as the new starting point would bias the optimizer toward nearby local minima, artificially smoothing the resulting front and making it impossible to assess whether the underlying trade-off is genuinely continuous or merely an artifact of the initialization strategy. As a consequence, a single run of the procedure displays small irregularities, and occasional local non-monotonicities, reflecting the fact that the global optimizer does not always converge to an equally good optimum for nearby target values; nonetheless, the overall trend is smooth and continuous, confirming that the trade-off between fidelity and success probability is a genuine feature of the physical system, and not an artifact of the optimization procedure. The full source code used for the simulations is provided in Appendix A.1, for completeness.

4.3.3 Multi-run aggregation

A single execution of the scheme described above already produces a full Pareto front, since the 80 target values of p_{target} are scanned within the same run, each optimization being carried out independently via its own MLSL global search. However, because the initial starting points sampled by MLSL in the global stage are drawn randomly

(via the low-discrepancy sequence), different executions of the same script can converge to slightly different local optima for the same p_{target} , depending on the random initialization, with the irregularities described above. To obtain a more robust and representative front for each value of Γ , the full optimization script is executed N independent times, each producing an independent set of 80 points $(\mathcal{F}^*, p_{\text{succ}}^*)$.

For each target probability, we then compare the corresponding results across all N runs and retain the one with the highest fidelity. The Pareto front defined in the plots reported throughout this Chapter is, therefore the *upper envelope*, in fidelity, over the N independent runs, rather than the output of any single execution; this is the source of the smooth, near-monotonic behaviour observed in most of the reported fronts, as opposed to the more irregular front that would be obtained from a single run alone.

In addition to this consolidated Pareto front, we also generate a set of supplementary physical plots for a selected configuration, by default the one achieving the highest fidelity across the entire dataset (unless a specific iteration is manually selected). These include the real and imaginary parts of the resulting gate $G_{4 \times 4}$, displayed as three-dimensional bar plots; the time evolution of the photon density across the six waveguides for each of the four logical input states, shown as a density map over the propagation length T ; and the corresponding photon density distribution at the final time $T = 1$, displayed as a bar plot for each input state. These plots are used in the following sections to provide a physical interpretation of the optimized gates beyond the aggregate figures of merit $\mathcal{F}$ and p_{succ} .

Again, the full source code used for the simulations is provided in Appendix A.2, for completeness.

4.4 Numerical results: CNOT in linear regime

Let us start by considering the linear case, i.e., $\Gamma = 0$, in which the self-Kerr term is absent and the Hamiltonian reduces to a standard tight-binding model

$$\mathcal{H} = \sum_{m=1}^6 E_m a_m^\dagger a_m + \sum_{m=1}^5 J_{m,m+1} (a_m^\dagger a_{m+1} + a_{m+1}^\dagger a_m). \quad (4.18)$$

In this regime, the system implements a linear-optical quantum walk, and the results of the optimization can be directly compared with two established reference points in this thesis.

The first reference result to be compared to is represented by the already mentioned *the no-go theorem* by Ralph et al. [10], which establishes that any probabilistic CNOT gate implemented with linear optics, dual-rail encoding, and without ancilla photons is bounded by a maximum success probability of $p_{\text{succ}} = 1/9$ at unit fidelity. This bound sets the theoretical ceiling for the linear quantum walk case as well, and the optimization should saturate it if the six-waveguide geometry is well enough designed to achieve optimality.

The second reference point is the optimized *photonic CNOT gate* already discussed in the previous Chapter [15], which demonstrated that precisely this bound can be achieved within a six-site waveguide lattice under quantum walk dynamics, with the same dual-rail encoding and the same Hamiltonian structure adopted here (see, e.g., Fig. 3.3). Since our theoretical setup in the linear limit is essentially identical to that of

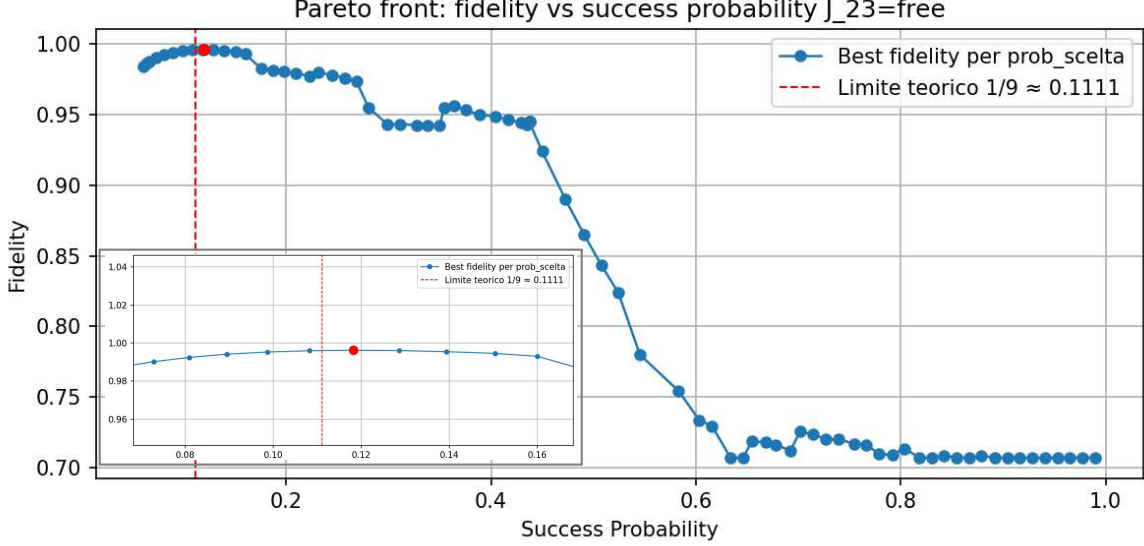


Figure 4.3: Optimized Pareto front obtained from the algorithmic optimization procedure in the linear regime (i.e., for $\Gamma = 0$). The inset shows a zoom around the theoretical bound $p_{\text{succ}} = 1/9$, where the fidelity reaches its maximum (close to 100%) value.

Lahini et al., we expect the optimization to reproduce their result, both in terms of the maximum success probability and in terms of the structure of the optimal parameters. The only difference from the setup of Lahini et al. lies in the number of free parameters. In their work, both the on-site energy of the first waveguide (i.e., E_1) and the hopping coefficient J_{23} between waveguides 2 and 3 are fixed to zero, reducing the free parameters from 10 to 9. The constraint $J_{23} = 0$ is deliberately imposed to replicate the traditional scheme of Ralph et al. [10], in which no coupling is allowed between the first two sites (i.e., the ancillary site $|a_1\rangle$ and the first site of the control qubit $|c_1\rangle$), and the other four sites.

In our optimization, instead, J_{23} is left as a free parameter: we decided not to deliberately enforce the architecture from Ralph et al., but rather ask whether the optimizer reproduces it spontaneously. As it will be shown in the results below, this is indeed the case: interestingly, the optimal value of J_{23} found by the optimizer is of the order 10^{-5} , i.e., $J_{23} \rightarrow 0$, thus confirming that the decoupling condition emerges naturally as the optimal configuration, without being imposed as a constraint.

So, following the optimization strategy outlined in Sec. 4.3, we have numerically obtained the Pareto front shown in Fig. 4.3, in which we report the best achieved fidelity as a function of the success probability in the linear regime. A few points worth being discussed are highlighted in the following paragraphs.

Maximal fidelity and optimal success probability

We can immediately notice a striking result from Fig. 4.3: the optimal configuration achieving the highest fidelity of all linear-regime runs corresponds to a target success probability of $p_{\text{target}} = 0.1182$, which is very close to the theoretical bound $1/9 \approx 0.1111$. All the optimized model parameters are given in units of the largest hopping coefficient of the optimized set of p_{target} , defined as $J_{\text{max}}^{\text{loc}} = \max_{ij} |J_{ij}^*|$. For this optimization,

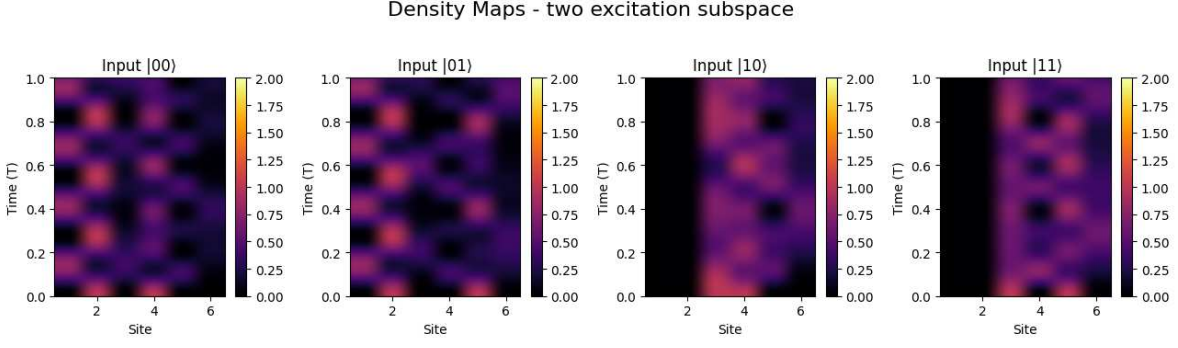


Figure 4.4: Intensity graph of the time evolution of the photon density across the six waveguides, from $T = 0$ to $T = 1$, obtained after initializing each computational basis state, at the best-fidelity point (corresponding to $p_{\text{succ}} \simeq 1/9$, $\mathcal{F} = 0.9962$), and for $\Gamma = 0$.

$J_{\text{max}}^{\text{loc}} = J_{12}^* \approx 3.272 \pi^{-1}$, and the optimized parameter vector thus reads

$$\begin{aligned} \frac{\boldsymbol{\theta}^*}{J_{\text{max}}^{\text{loc}}} &= \left(\frac{E_2}{J_{\text{max}}^{\text{loc}}}, \dots, \frac{E_6}{J_{\text{max}}^{\text{loc}}}, \frac{J_{12}}{J_{\text{max}}^{\text{loc}}}, \dots, \frac{J_{56}}{J_{\text{max}}^{\text{loc}}} \right) = \\ &= (0.975, -0.385, -0.840, 0.650, 0.580, \\ &\quad 1.000, 1.4 \cdot 10^{-4}, 0.719, 0.940, 0.498). \end{aligned} \quad (4.19)$$

We notice that the second hopping coefficient, $J_{23}/J_{\text{max}}^{\text{loc}} = 1.4 \cdot 10^{-4}$, is numerically close to zero, thus confirming that the optimizer spontaneously reproduces the decoupling condition imposed by hand in Lahini et al. The resulting figures of merit in the optimal configuration corresponding to the array above are

$$p_{\text{succ}} = 0.11822, \quad \mathcal{F} = 0.99623, \quad (4.20)$$

with a residual cost $C(\boldsymbol{\theta}^*) = 7.52 \cdot 10^{-3}$. The corresponding gate matrix, projected onto the logical subspace, is

$$G_{4 \times 4} = \begin{pmatrix} 0.377 - 0.010i & 0.015 + 0.003i & 0 & 0 \\ 0.015 + 0.003i & 0.303 - 0.033i & 0 & 0 \\ 0 & 0 & 0.001 + 0.007i & 0.344 - 0.022i \\ 0 & 0 & 0.344 - 0.022i & 0.007 - 0.001i \end{pmatrix}, \quad (4.21)$$

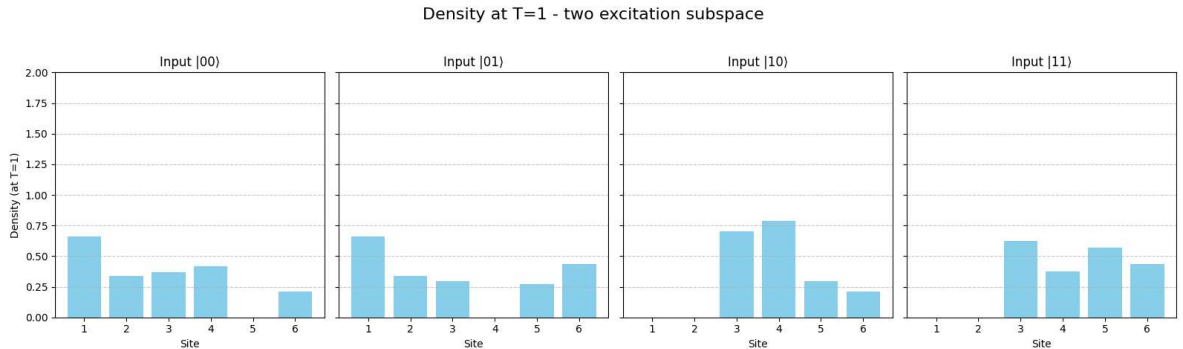


Figure 4.5: Photon density at $T = 1$ across the six waveguides, for each computational basis input state, for parameters as in Fig. 4.4.

which compares well with the one reported in Fig. 3.4(a). Its normalized, phase-corrected version reads

$$\tilde{G} = \begin{pmatrix} 0.548 & 0.022 & 0 & 0 \\ 0.022 & 0.443 & 0 & 0 \\ 0 & 0 & 0.010 & 0.501 \\ 0 & 0 & 0.501 & 0.010 \end{pmatrix}, \quad (4.22)$$

where the entries below 10^{-3} have been set to zero for readability. The block structure of $\tilde{G}$ is then evident: the upper-left 2×2 block acts on the $\{|00\rangle, |01\rangle\}$ subspace (control qubit in state $|0\rangle$), while the lower-right block acts on $\{|10\rangle, |11\rangle\}$ (control qubit in state $|1\rangle$), with small off-diagonal coupling between the two subspaces. Within the second block, the dominant off-diagonal entries implement the required bit-flip on the target qubit, in agreement with the expected CNOT action.

It is worth noting that the entries of $\tilde{G}$ should not be compared directly with those of U_{CNOT} , whose non-zero entries are all equal to 1. In fact, since $\tilde{G}$ is normalized to the unit Frobenius norm by construction, the correct reference for comparison is the normalized CNOT matrix, i.e., $U_{\text{CNOT}}/\|U_{\text{CNOT}}\|_F = U_{\text{CNOT}}/2$, whose non-zero entries are all equal to $1/2$. The dominant entries of $\tilde{G}$ (0.501 and 0.548, 0.443) are close to this value, and the small deviations from $1/2$ are consistent with the fidelity $\mathcal{F} = 0.99623$. To complement the discussion above, we show the raw temporal evolution of the photon density together with the projected logical gate, both evaluated at the best-fidelity point. In particular, Fig. 4.4 shows the time evolution of the photon density across the six sites, from $T = 0$ to $T = 1$, for each of the four computational basis input states, using the optimized Hamiltonian parameters. Fig. 4.5 shows the same evolution evaluated only at the final time $T = 1$, as a bar graph on the six sites. Since this is the raw, un-projected output of the evolution, it does not by itself resemble a CNOT gate: a sizable fraction of the population remains outside the logical subspace, consistent with the modest success probability (i.e., only slightly larger than 10%) at this point of the Pareto front.

By contrast, Fig. 4.6 shows the gate $G_{4 \times 4}$ obtained after projecting the evolution onto the logical subspace, normalizing, and correcting for the global phase: $\tilde{G}$. In this representation, the CNOT structure is immediately recognized: the real part reproduces the identity action on $|00\rangle$ and $|01\rangle$, and the swap between $|10\rangle$ and $|11\rangle$, while the imaginary part remains close to zero throughout.

The tail of the Pareto front at larger success probability

As a representative example of the behaviour in the high-probability tail of the Pareto front, we report the optimized configuration for $p_{\text{target}} = 0.9409$, where the fidelity has dropped to $\mathcal{F} = 1/\sqrt{2} \approx 0.7071$. Expressing all parameters in units of $J_{\text{max}}^{\text{loc}} = |J_{45}^*| \approx 2.740 \pi^{-1}$, the optimized parameter vector now reads

$$\frac{\boldsymbol{\theta}^*}{J_{\text{max}}^{\text{loc}}} = (-1.201, 1.356, -0.251, -0.251, -0.148, -0.520, 1.5 \cdot 10^{-5}, -0.012, -1.000, 3.8 \cdot 10^{-3}), \quad (4.23)$$

in which, as before, $J_{\text{max}}^{\text{loc}}$ is the largest hopping coefficient in absolute value for this configuration. In particular, $J_{23}/J_{\text{max}}^{\text{loc}} \approx 1.5 \cdot 10^{-5}$ and $J_{56}/J_{\text{max}}^{\text{loc}} \approx 3.8 \cdot 10^{-3}$ are negligible,

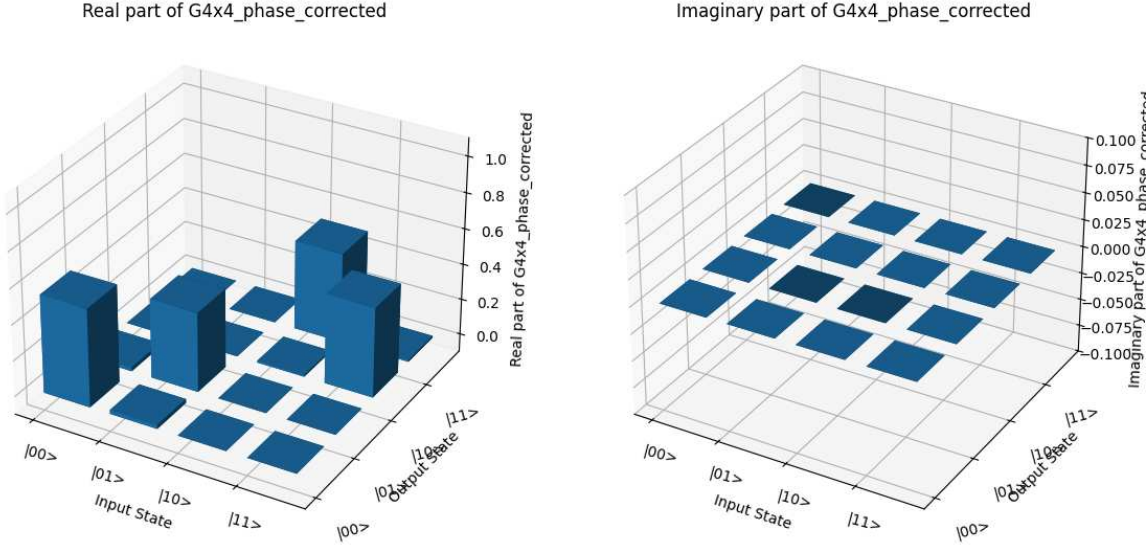


Figure 4.6: Three-dimensional histogram of the gate matrix, projected onto the logical subspace, normalized and de-phased, $\tilde{G}$, for $\Gamma = 0$.

and $J_{34}/J_{\max}^{\text{loc}} \approx -0.012$ is also essentially zero, thus producing a decoupling of the two qubits. Only J_{12} and J_{45} remain as active hopping channels, respectively. In fact, the hopping between the first ancilla and the first site of the control qubit, and the hopping between the two sites of the target qubit are preserved. The figures of merit resulting from the algorithmic optimization are

$$p_{\text{succ}} = 0.9409, \quad \mathcal{F} = \frac{1}{\sqrt{2}} \approx 0.7071, \quad (4.24)$$

with a residual cost $C(\boldsymbol{\theta}^*) = 0.5000$, which is consistent with $1 - \mathcal{F}^2 = 1/2$ at the penalty minimum.

The consequences of this parameters configuration are directly visible in the time evolution of the photon density, as shown in Fig. 4.7, and in the corresponding density at $T = 1$, as shown in Fig. 4.8, for each computational basis input state. For inputs $|00\rangle$ and $|01\rangle$ (control qubit in $|c_1\rangle$), the population oscillates only between sites 1 and 2, with no leakage onto site 3; conversely, for inputs $|10\rangle$ and $|11\rangle$ (control qubit in $|c_2\rangle$),

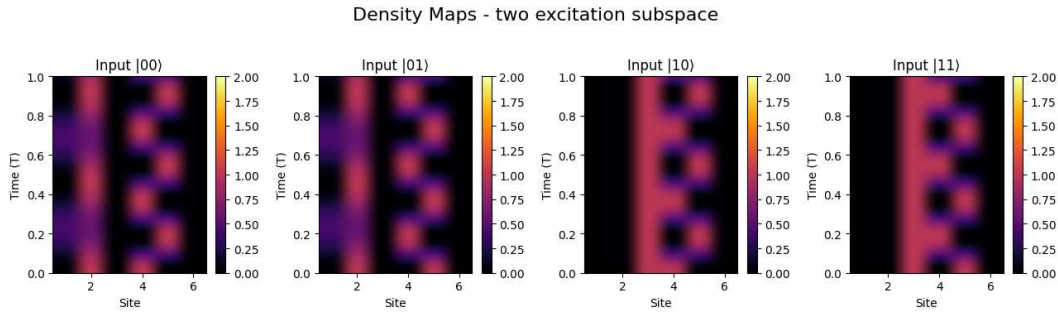


Figure 4.7: Time evolution of the photon density across the six waveguides, from $T = 0$ to $T = 1$, for each computational basis input state, at a representative tail point on the Pareto front (i.e., $p_{\text{succ}} = 0.9409$, $\mathcal{F} = 1/\sqrt{2}$), for $\Gamma = 0$.

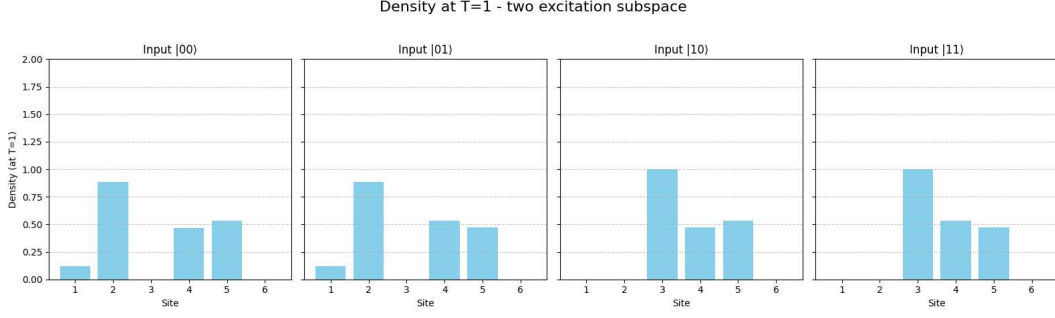


Figure 4.8: Photon density at $T = 1$ across the six waveguides, for each computational basis input state, at a representative tail point on the Pareto front, as in Fig. 4.7.

the population remains pinned to site 3 throughout the evolution, with unit density and no oscillation at all. This is the direct signature of $J_{23} \approx 0$: the control-qubit photon is confined to whichever of the two control sites it started in, and cannot access the other. The target block (sites 4 and 5), by contrast, oscillates in the same way regardless of the control input, confirming that its evolution is independent of the control qubit’s state (due to $J_{34} \approx 0$).

The optimized matrix, after projecting onto the logical subspace, normalizing, and de-phasing, explicitly obtained as

$$\tilde{G} = \begin{pmatrix} 0.332 & 0.353 & 0 & 0 \\ 0.353 & 0.332 & 0 & 0 \\ 0 & 0 & 0.353 & 0.376 \\ 0 & 0 & 0.376 & 0.353 \end{pmatrix}, \quad (4.25)$$

whose entries below 10^{-3} have been set to zero, and is shown in Fig. 4.9. The two-block structure is preserved, confirming that the decoupling due to $J_{23} \rightarrow 0$ persists even at high success probability. However, the entries are now far from the reference values of $1/2$ expected for a perfect normalized CNOT, reflecting the significant loss of fidelity: in particular, the diagonal and off-diagonal entries within each block are comparable in magnitude, indicating that the gate is no longer implementing a clean bit-flip operation on the target qubit.

Now, going back to analyse the whole Pareto front in Fig. 4.3, it is clear that beyond $p_{\text{succ}} \approx 0.44$, the optimizer begins to show signs of difficulty in satisfying the target probability constraint: the achieved success probability stagnates for several consecutive target values, suggesting that the high-fidelity family of solutions has reached the boundary of its accessible region in the fidelity–success plane. As the target is further pushed, the optimizer transitions to a qualitatively different family of solutions, and the fidelity drops clearly, stabilizing at $\mathcal{F} = 1/\sqrt{2} \approx 0.707$ for all $p_{\text{succ}} \gtrsim 0.63$.

Inspection of the optimized parameters in this region shows that not only J_{23} but also J_{34} tends to zero, progressively decoupling the two qubits; correspondingly, the optimizer converges to configurations belonging to multiple distinct families, all sharing the same fidelity ceiling. This behaviour is qualitatively consistent with the existence of competing families of solutions with low fidelity, as reported, e.g., by Smith et al. [73] for the linear-optical CNOT gate. Since this region corresponds to a regime of low CNOT gate quality, thus of limited practical interest, a detailed characterization of these solution families is left as a possible direction for future investigation.

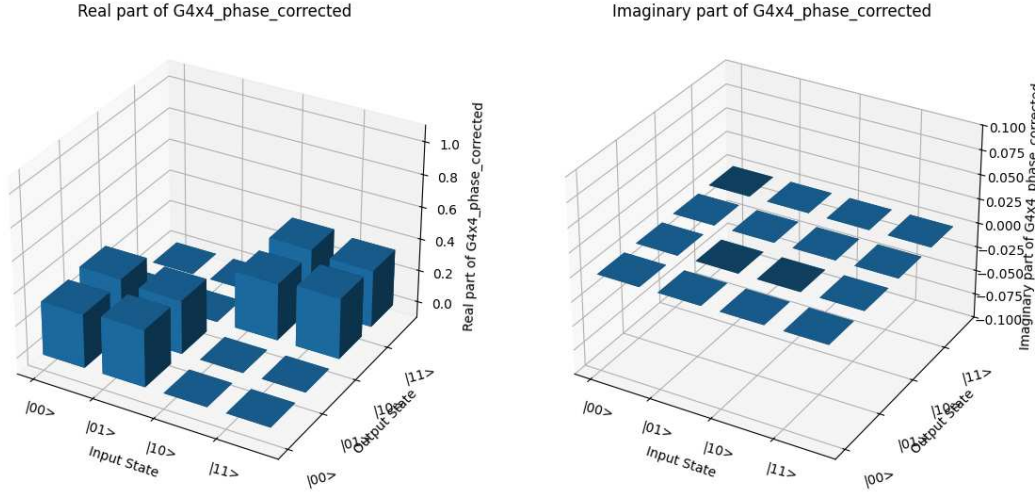


Figure 4.9: Real and imaginary parts of the normalized, de-phased logical gate $\tilde{G}$ at a representative tail point on the Pareto front (i.e., $p_{\text{succ}} = 0.9409$, $\mathcal{F} = 1/\sqrt{2}$), showing the factorized, product-like block structure discussed in the text, for $\Gamma = 0$.

Comparison with the solution from Lahini et al.

As a final comparison, we compare our results for the linear quantum walk case with the ones reported by Lahini et al. [15] and already discussed in Fig. 3.3, shown side by side in Fig. 4.10 for easier visualization. As it is immediately evident, the two curves agree qualitatively in the high-fidelity, low-probability region: both reach fidelities close to unity around $p_{\text{succ}} \approx 1/9$, consistent with the known theoretical bound.

However, in the tail the two results differ noticeably: our data saturate sharply at $\mathcal{F} = 1/\sqrt{2} \approx 0.7071$ for $p_{\text{succ}} \gtrsim 0.6$, whereas the fidelity reported by Lahini et al. decreases monotonically as a function of p_{succ} , without an analogous plateau and reaching values as low as $\mathcal{F} \approx 0.25$ at $p_{\text{succ}} = 1$.

Interestingly, an independent piece of evidence for the same physical floor comes from Scala et al. [17], in which a distinct architecture is optimized (as already mentioned in the previous Chapter) by using a different optimizer. In Ref. [17], the average gate fidelity was defined as the mean of the squared diagonal overlaps, i.e., $\bar{F} = |S|^{-1} \sum_i |\langle i | \tilde{U}^\dagger T | i \rangle|^2$, rather than the modulus of the average overlap used in this work. For the noiseless linear case (i.e., $U = 0$ in the parametrization of nonlinearity given in that work), average gate fidelity plateaus at $\bar{F} = 0.5$ are obtained regardless of the number of blocks. Since a product-gate approximation gives $\bar{F} = \mathcal{F}^2$ under this alternative convention, the result reported there is fully consistent with our $\mathcal{F} = 1/\sqrt{2}$ floor, although neither work explicitly comment on this feature. We regard this agreement, obtained with an unrelated optimization scheme and circuit parametrization, as independent support for the physical origin of the tail plateau discussed in this Chapter.

To conclude, it is fair to say that the origin of the discrepancy with the tail behaviour reported by Lahini et al. remains unclear to us, at this stage. No supplementary material or code accompanying the work in Ref. [15] is publicly available, and repeated verification of our implementation against their formalism did not reveal any discrepancy. We initially conjectured that the difference might stem from their choice of fixing $J_{23} = 0$ by construction throughout the whole Pareto front, whereas in this work J_{23} is left free

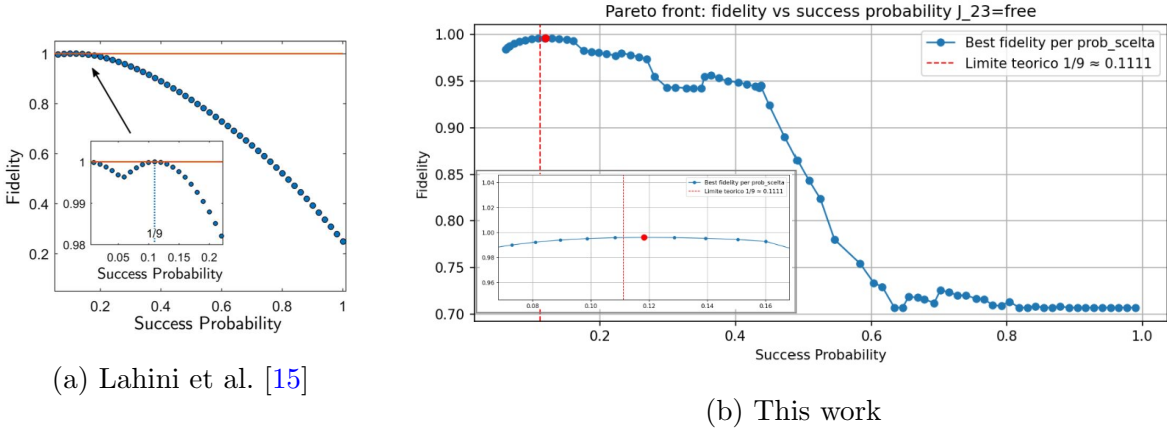


Figure 4.10: Direct comparison between (a) the fidelity-vs-success-probability Pareto front reported by Lahini et al. and (b) the one obtained in this work for the linear case, $\Gamma = 0$. Both show near-unit fidelity around the theoretical bound $p_{\text{succ}} \approx 1/9$, but differ substantially in the tail behaviour, as discussed in the text.

and only vanishes spontaneously in the tail. However, repeating our optimization with J_{23} fixed to zero as an initial condition yields the same $\mathcal{F} = 1/\sqrt{2}$ plateau in the tail, ruling out this explanation: the floor appears to be a robust consequence of the fact that once the control and target qubits are fully decoupled, each evolves independently of the other, and no such independent evolution can reproduce an entangling operation like the CNOT beyond $\mathcal{F} = 1/\sqrt{2}$, regardless of whether J_{23} is optimized freely or fixed to 0. The reason for the qualitatively different, monotonically decreasing tail reported by Lahini et al. remains therefore an open question.

4.5 CNOT gate in nonlinear quantum walks

We now turn to the case $\Gamma \neq 0$, in which the self-Kerr term of Eq. (4.1) is switched back on, and the full Hamiltonian

$$\mathcal{H} = \sum_{m=1}^6 E_m a_m^\dagger a_m + \sum_{m=1}^5 J_{m,m+1} (a_m^\dagger a_{m+1} + a_{m+1}^\dagger a_m) + \Gamma \sum_{m=1}^6 a_m^\dagger a_m (a_m^\dagger a_m - 1) \quad (4.26)$$

is used to generate the evolution. This solution has not been provided previously in the literature, and it thus constitutes a genuine original contribution of the present thesis. Physically, the nonlinearity introduces an energy cost whenever two photons simultaneously occupy the same waveguide, which occurs during the intermediate propagation even though the computational basis states of Eq. (4.4) always contain exactly one photon per qubit. As discussed in the linear case, the bound at $p_{\text{succ}} = 1/9$ found by Ralph et al. is a direct consequence of the fact that linear optics alone cannot induce the kind of conditional, photon-photon interaction required by an entangling gate; the self-Kerr term is precisely the ingredient that reintroduces such an interaction, and the central question of this section is: to what extent, and at which cost in nonlinearity strength, this allows the optimization to move beyond the linear bound?

Following the convention adopted by Scala et al. [17], we express the strength of the nonlinearity relative to the natural energy scale of the hopping, J_{max} , rather than in

Γ	$E_{\max} = J_{\max}$	$\Gamma/J_{\max}$
0	4	0
0.01	4.5	0.0022
0.1	5	0.020
0.25	5.5	0.045
0.5	6	0.083
0.75	7	0.107
1	8	0.125

Table 4.1: Optimization bounds $E_{\max}$ used for each value of Γ , and the corresponding normalized nonlinearity $\Gamma/J_{\max}$.

absolute units of Γ alone. In our case, $J_{\max}$ is not extracted a posteriori from the optimized parameters, but coincides with the bound $E_{\max}$ of the symmetric search interval $[-E_{\max}, E_{\max}]$ used for the optimization at each value of Γ , as discussed in Section 4.1: since this bound must be widened as Γ increases to avoid artificially constraining the optimizer, it grows together with Γ itself, and the two must be considered jointly. The values used throughout this work are summarized in Table 4.1.

Since $E_{\max} = J_{\max}$ is increased together with Γ , the normalized ratio $\Gamma/J_{\max}$ grows more slowly than Γ itself, reaching only $\Gamma/J_{\max} \approx 0.125$ at $\Gamma = 1$. This is worth contrasting with the regime identified by Scala et al.[17], who find their best CNOT performance around $U/J_{\max} \approx 0.5$, with a disadvantageous photon-blockade effect setting in only for $U/J_{\max} \gtrsim 1$. All the values of Γ explored in this work therefore remain, in this normalized sense, in a comparatively weak-nonlinearity regime relative to the hopping strength. We will discuss this point more once the results of this section have been presented.

Essentially, we repeated the optimization procedure described in Sec. 4.3 for a broad range of Γ values, starting from vanishingly small nonlinearities down to $\Gamma/J_{\max} \sim 10^{-4}$, up to $\Gamma/J_{\max} = 0.125$. For $\Gamma/J_{\max} \lesssim 0.02$, the resulting front is essentially indistinguishable from the linear case discussed above, both in the position of the high-fidelity peak as well as in the shape of the Pareto front tail. A qualitatively different behaviour first becomes visible around $\Gamma/J_{\max} \approx 0.045$, and it becomes increasingly pronounced as the nonlinearity is further increased.

4.5.1 Weak non-linearities

In Fig. 4.11 we immediately report the main result of this work, showing the optimized Pareto fronts for weakly nonlinear quantum walks. In the following, we will focus on selected results in order to give more extensive discussions. First, let us notice that the linear result is here reported for direct comparison, which allows to confirm the picture built up progressively in the previous sections: as Γ increases, the front does not simply shift upward uniformly, but undergoes a qualitative reorganization. At $\Gamma = 0$, a single family of solutions gradually decrease from the Ralph bound down to the floor $\mathcal{F} = 1/\sqrt{2}$. At intermediate Γ (0.5, 0.75), this single-family picture breaks down: the $J_{23} \approx 0$ family inherited from the linear regime now competes with a second, structurally different one in which J_{12} , J_{23} , and J_{56} all vanish and only the three-site

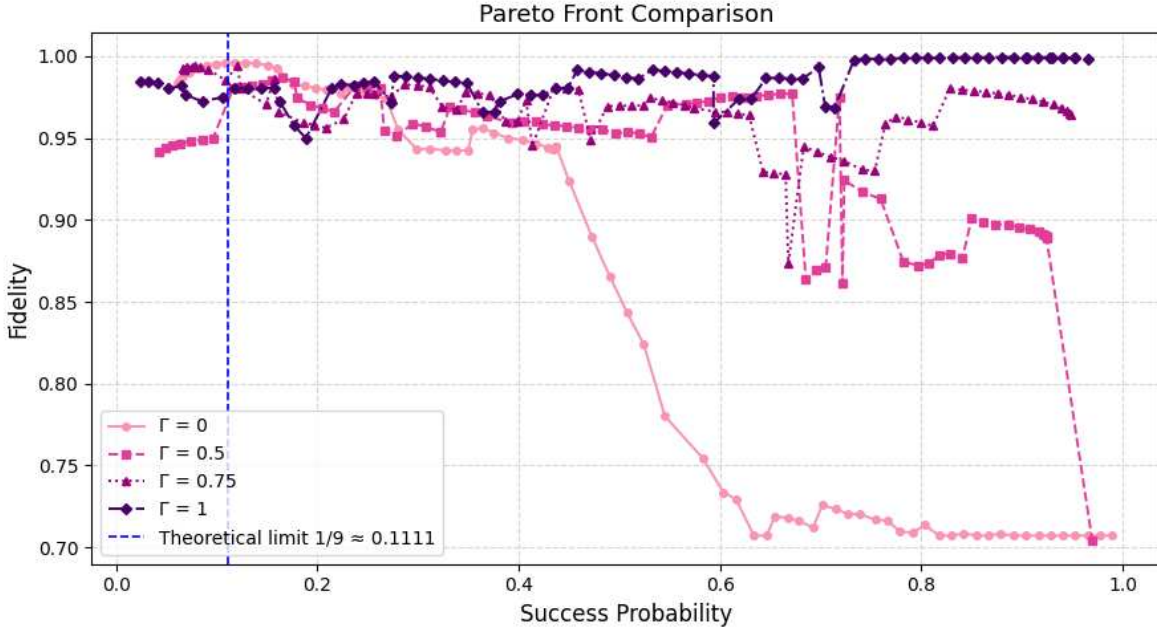


Figure 4.11: Numerically optimized Pareto fronts for the cases of linear or weakly nonlinear quantum walks, i.e., $\Gamma = 0, 0.5, 0.75, 1$.

block $J_{34}-J_{45}$ remains active, with neither family clearly dominating over the full range of p_{succ} ; this unresolved competition is what produces the irregular, non-monotonic fronts observed above. As we will discuss in detail below, starting from $\Gamma = 1$ this competition resolves in favour of a single, robust family capable of sustaining near-unit fidelity across virtually the entire range of success probabilities, up to unity (i.e., deterministic gating).

Let us now focus our discussion on the two regions of greatest physical interest: the neighbourhood of the linear-optical bound around $p_{\text{succ}} = 1/9$, and the near-deterministic region $p_{\text{succ}} \rightarrow 1$, respectively.

Near the LOQC bound, all the values of Γ perform comparably well, with $\Gamma = 0$ actually achieving the single highest fidelity of the whole dataset ($\mathcal{F} = 0.9962$); the nonlinearity does not provide any appreciable advantage here.

The picture is entirely different in the near-deterministic region: here, $\Gamma = 1$ is the only value maintaining $\mathcal{F} > 0.997$ for $p_{\text{succ}} \gtrsim 0.83$, while every other value of Γ remains capped well below $\mathcal{F} = 0.9$ (while for $\Gamma = 0$ the Fidelity collapses to the $1/\sqrt{2}$ floor). This confirms that the principal benefit of the self-Kerr nonlinearity investigated in this work is not an improvement over the linear optimum, but the ability to sustain high fidelity at high success probability, up to and including deterministic operation. We consider the latter result one of the main outcomes of our work.

4.5.2 Strong non-linearities

As a qualitative exploration beyond the regime studied in Fig. 4.11, we report in Fig. 4.12 the Pareto fronts obtained from single optimization runs when larger values of the photon-photon nonlinearity are assumed, i.e., $\Gamma \in \{10, 20, 30\}$. These runs were performed with the same computational resources as in the systematic multi-run anal-

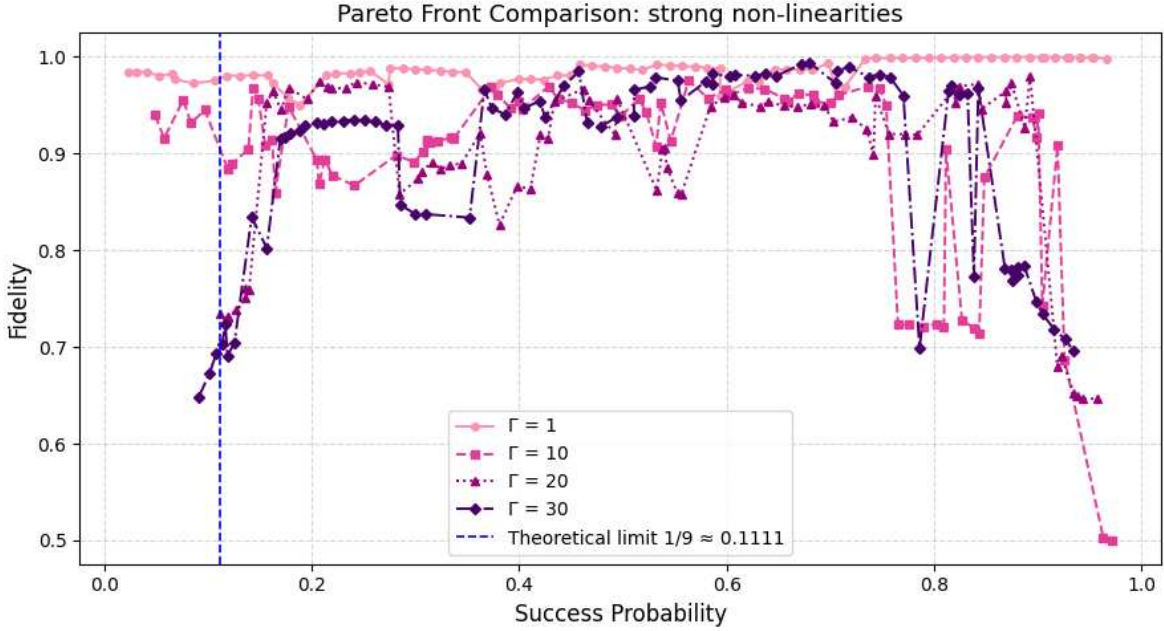


Figure 4.12: Numerically optimized Pareto fronts for the cases of strongly nonlinear quantum walks, i.e., $\Gamma \geq 10$. The results for $\Gamma = 1$, already reported in the previous Figure, are here shown for direct comparison.

ysis of $\Gamma = 1$ described below. However, these are insufficient for this case, which would require a much broader bound, which implies a greater number of iterations to explore it with a gradient free optimizer, and therefore more computational power. Therefore, these runs should be regarded as indicative rather than convergent results. For direct comparison, we report again in this Figure the same results already shown for $\Gamma = 1$ in Fig. 4.11. Nevertheless, the overall trend is clear: for large Γ , the fidelity–success trade–off becomes increasingly irregular and the front fails to maintain high fidelity across the full range of success probabilities. This is consistent with the findings of Scala et al. [17], who observed that excessively large nonlinearity, pushing the system toward the photon blockade regime, is detrimental to gate performance, and that optimal operation occurs for moderate nonlinearity values (i.e., $\Gamma/J_{\max} \lesssim 0.5$). A thorough characterization of this large- Γ regime, requiring significantly greater computational resources, is left for future work. We further notice, in this respect, that this nonlinearity regime might be experimentally unattainable in realistic material platforms, and its investigation could be regarded as a purely academic interest.

In what follows, we will focus the discussion on the most significant result, i.e. the Pareto front obtained for $\Gamma = 1$, for which the deviation from the linear behaviour is most significant; the remaining intermediate cases are reported in Appendix B for completeness.

4.5.3 Focus on nonlinear results for $\Gamma = 1$ ($\Gamma/J_{\max} = 0.125$)

In Fig. 4.13 we isolate the optimized Pareto front for $\Gamma = 1$, i.e., the nonlinearity that gives the most striking result reported so far in this work. For comparison with Scala et al. [17], we notice that this value corresponds to the normalized $\Gamma/J_{\max} = 0.125$ for this

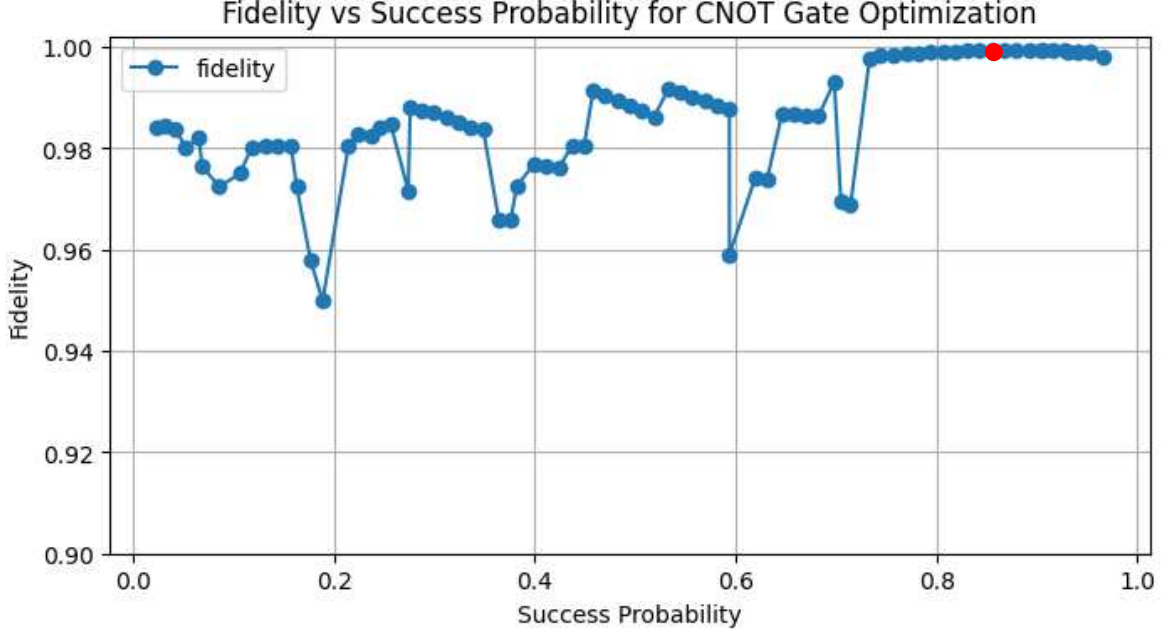


Figure 4.13: Optimized Pareto front for $\Gamma = 1$ (i.e., $\Gamma/J_{\max} = 0.125$), here reported again for close inspection.

batch of simulations. Unlike the other cases shown in Fig. 4.11, the front is remarkably smooth and uniformly high until very large success probability. In fact, the Fidelity remains above $\mathcal{F} \approx 0.95$ throughout the entire range of $0 < p_{\text{succ}} < 1$, with only two modest, isolated dips, around $p_{\text{succ}} \approx 0.59$ ($\mathcal{F} \approx 0.959$) and $p_{\text{succ}} \approx 0.71$ ($\mathcal{F} \approx 0.969$). Beyond $p_{\text{succ}} \approx 0.73$, the Pareto front remarkably stabilizes above $\mathcal{F} = 0.997$, and there remains all the way to $p_{\text{succ}} \rightarrow 1$, as shown in Fig. 4.14.

Inspecting the optimized parameters across the whole front shows a qualitative difference from the cases discussed previously: rather than a competition between multiple, comparably good families, $\Gamma = 1$ is dominated almost everywhere by a single robust family, with $J_{23} \approx 0$, decoupling the first two sites from the others, and, for much of the front, $J_{56} \approx 0$ as well, decoupling the last ancilla. The two isolated dips mentioned above are the only points at which this family appears to be locally interrupted, without being replaced by a persistent alternative family. Beyond these two dips, the same family persists uninterrupted into the near-deterministic regime, with fidelity increasing monotonically towards unity.

This confirms the trend anticipated in the previous sections: as Γ increases, the front does not merely flatten in an average sense, but converges towards a single, increasingly robust family of solutions capable of sustaining near-unit fidelity across the entire range of success probabilities. This behaviour does not hold forever, as shown in Fig. 4.12: when further increasing Γ , convergence towards high Fidelities remains problematic on approaching unitary success probability.

Highest-fidelity point The highest-fidelity configuration on this Pareto front is found for $p_{\text{target}} = 0.8795$, deep in the near-deterministic region of the optimized CNOT gate. It is worth noting that this point is only nominally the maximum: the fidelity remains essentially unchanged, within $5 \cdot 10^{-4}$, for all of the subsequent points up to

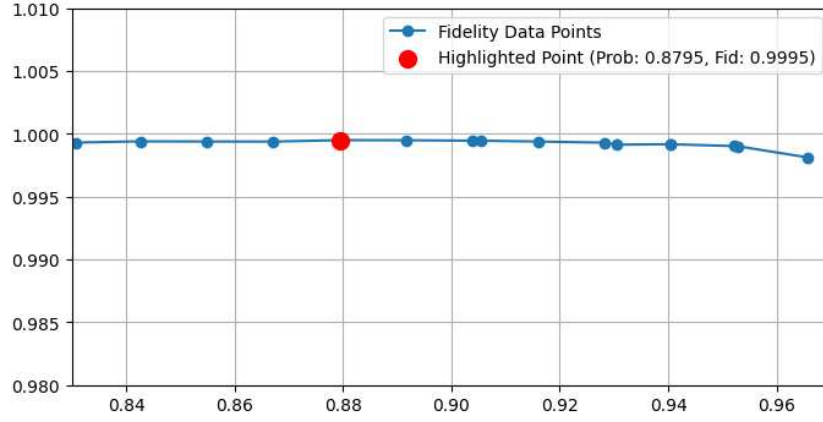


Figure 4.14: Zoom-in of the Pareto front shown in Fig. 4.13 for $\Gamma = 1$, in the near-deterministic region of the CNOT gate, i.e., $p_{\text{succ}} \gtrsim 0.83$. The highlighted (red) point corresponds to the highest-fidelity configuration, as discussed below.

$p_{\text{succ}} \rightarrow 1$ (see Fig. 4.14), such that this configuration should be understood as representative of the entire near-deterministic plateau, rather than as a uniquely optimal point. Expressing all the optimal parameters in units of $J_{\text{max}}^{\text{loc}} = |J_{45}^*| \approx 5.869 \pi^{-1}$, the optimized parameter vector reads

$$\frac{\boldsymbol{\theta}^*}{J_{\text{max}}^{\text{loc}}} = (0.756, -0.058, -0.788, -0.815, -1.007, -0.149, 1.8 \cdot 10^{-7}, -0.148, 1.000, -1.7 \cdot 10^{-5}), \quad (4.27)$$

Both $J_{23}/J_{\text{max}}^{\text{loc}} \approx 1.8 \cdot 10^{-7}$ and $J_{56}/J_{\text{max}}^{\text{loc}} \approx -1.7 \cdot 10^{-5}$ are negligible, confirming the decoupling mechanism found throughout this chapter, that persists even in this strongly non-linear near-deterministic regime.

The resulting figures of merit are

$$p_{\text{succ}} = 0.8795, \quad \mathcal{F} = 0.99950, \quad (4.28)$$

with a residual cost $C(\boldsymbol{\theta}^*) = 1.00 \cdot 10^{-3}$, the smallest of all the configurations discussed in this chapter. The corresponding gate matrix, projected onto the logical subspace, is

$$G_{4 \times 4} = \begin{pmatrix} 0.937 + 0.029i & -0.001 - 0.011i & 0 & 0 \\ -0.001 - 0.011i & 0.937 + 0.021i & 0 & 0 \\ 0 & 0 & -0.024 + 0.041i & 0.937 + 0.025i \\ 0 & 0 & 0.937 + 0.025i & -0.024 - 0.021i \end{pmatrix}, \quad (4.29)$$

in which the elements significantly different from zero tend now to be close to 1, as a further indication that the CNOT gate tends to be deterministic. Its normalized, phase-corrected version reads

$$\tilde{G} = \begin{pmatrix} 0.500 & 0.006 & 0 & 0 \\ 0.006 & 0.500 & 0 & 0 \\ 0 & 0 & 0.025 & 0.500 \\ 0 & 0 & 0.500 & 0.017 \end{pmatrix}, \quad (4.30)$$

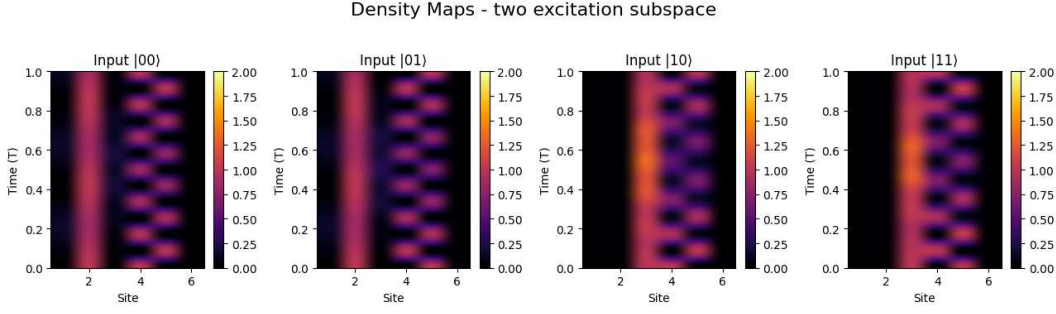


Figure 4.15: Time evolution of the photon density across the six waveguides, from $T = 0$ to $T = 1$, for each computational basis input state, at the highest-fidelity point for $\Gamma = 1$.

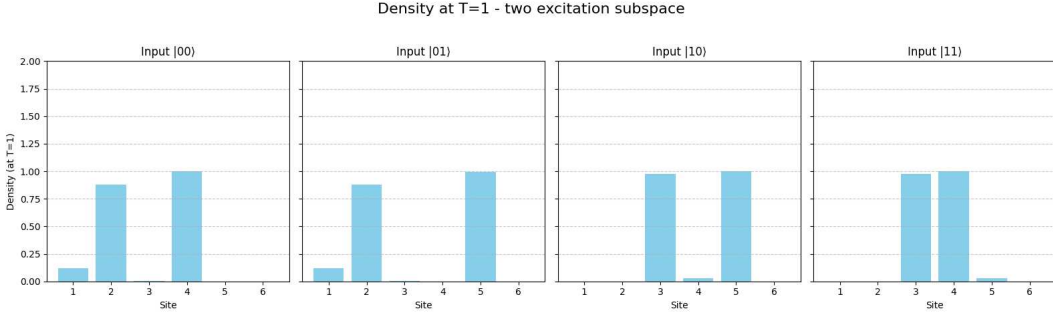


Figure 4.16: Photon density at $T = 1$ across the six waveguides, for each computational basis input state, at the highest-fidelity point for $\Gamma = 1$.

in which again the entries with modulus below 10^{-3} have been set to zero. The dominant entries are now essentially indistinguishable from the ideal value $1/2$ expected for a *perfect normalized CNOT*, with deviations at the 10^{-3} level. This is a marked improvement over all other configurations presented in this chapter and it is consistent with the near-unit fidelity value obtained from the optimization $\mathcal{F} = 0.99950$.

Figures 4.15 and 4.16 show the time evolution of the photon density, and its value at $T = 1$, respectively, for each computational basis input state. The ancillary site $|a_2\rangle$ (i.e., site 6) remains essentially unpopulated throughout the evolution for all of the four inputs, which is consistent with the optimal value $J_{56} \approx 0$: in this near-deterministic regime, the second ancilla basically plays no dynamical role. The other ancillary site, $|a_1\rangle$ (i.e., site 1), does retain a small residual population (approximately 11%) for the inputs with the control qubit in $|c_1\rangle$ (sites $|00\rangle$, $|01\rangle$), consistent with the non-negligible value of J_{12} ; this residual loss is precisely what keeps $p_{\text{succ}} = 0.8795$ below unity, rather than reaching full determinism. For the remaining inputs, the density is essentially fully recombined onto the correct logical output sites at $T = 1$, with only percent-level residual leakage elsewhere. We notice the striking difference with the corresponding result shown in Fig. 4.5.

Finally, in Fig. 4.17 we report the corresponding projected, normalized, de-phased matrix $\tilde{G}$, whose near-uniform bar heights visually confirm the closeness to an ideal CNOT.

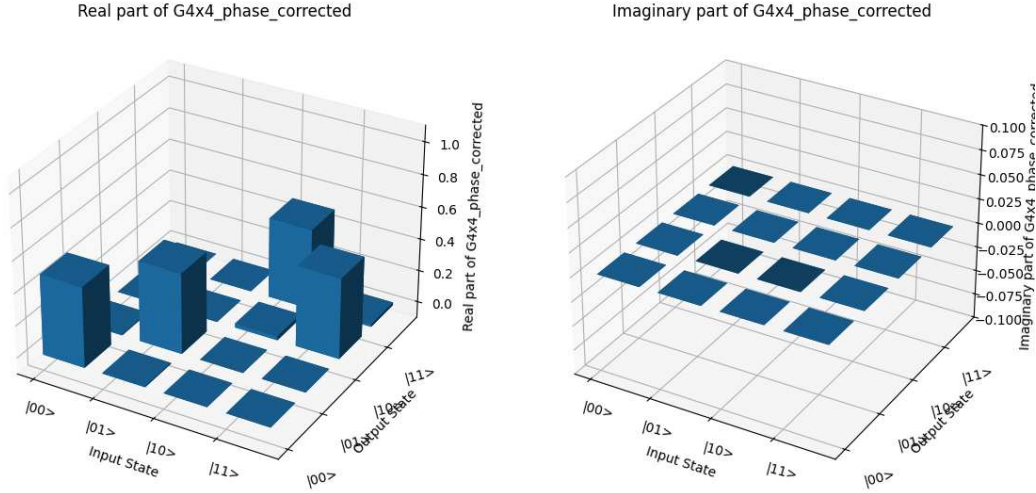


Figure 4.17: Real and imaginary parts of the normalized, dephased logical gate, $\tilde{G}$ at the highest-fidelity point for $\Gamma = 1$.

4.6 Physical plausibility of optimized parameters

A natural question is whether the optimized parameters found throughout this Chapter actually correspond to physically achievable couplings and nonlinearities in a physically realistic implementation, rather than to purely formal solutions of the optimization problem. While a definitive answer would require a specific target platform and fabrication constraints, an order-of-magnitude estimate can be obtained by fixing the dimensionless model parameters to the physical scale already introduced in Section 4.1: with $T_{\text{phys}} = L/v_g \approx 2$ ps and the condition $JT_{\text{phys}} \sim \pi$ yielding $\hbar J \approx 1$ meV for a model coupling of order unity, any of our optimized parameter J_{model} corresponds to a physical coupling of

$$\hbar J_{\text{phys}} \approx J_{\text{model}} \cdot 1 \text{ meV}. \quad (4.31)$$

Since Γ enters the Hamiltonian on exactly the same footing as the hopping terms, the same calibration applies to it. Under this estimate, the optimized $J_{\text{max}}^{\text{loc}}$ values found in this chapter (ranging from $J_{\text{max}}^{\text{loc}} \approx 3.3$ in the linear case to $J_{\text{max}}^{\text{loc}} \approx 6.8$ at $\Gamma = 0.75$) correspond to physical couplings of a few meV, somewhat larger than the ~ 1 meV reference value but not unreasonable, given that J depends exponentially on the physical separation between waveguides and can be increased simply by bringing them closer together.

A more informative comparison is obtained by considering the ratio Γ/J_{max} rather than the absolute values, since this ratio is independent of the specific calibration constant used above (both Γ and J scale with the same conversion factor, which therefore cancels). Scala et al. [17] report that realistic to optimistic values of the self-Kerr nonlinearity achievable in exciton-polariton waveguide platforms, given $\hbar J_{\text{max}} \sim 1$ meV, correspond to $U/J_{\text{max}} \approx 0.01$ – 0.5 . The values of Γ/J_{max} explored in this work (Table 4.1) range from $2.2 \cdot 10^{-3}$ to 0.125 , closely matching this physically motivated range: the bulk of this chapter's analysis ($\Gamma \leq 0.75$, i.e. $\Gamma/J_{\text{max}} \leq 0.125$) falls within of values considered achievable by Scala et al. This agreement, obtained without any deliberate tuning towards numbers from Scala et al., lends some physical credibility to the actual realization of the nonlinear regime explored in this chapter, although an objective

evaluation would require committing to a specific material platform and repeating the analysis of Section 4.1 with platform-specific values of L and v_g .

Summary

This chapter has explored whether the architectural simplicity of the linear-optical quantum-walk CNOT of Lahini et al. can be combined with a distributed self-Kerr nonlinearity, with the aim of overcoming the fundamental $p_{\text{succ}} = 1/9$ bound established from the theory of LOQC, while retaining the minimal six-waveguide geometry. In the linear case ($\Gamma = 0$), the optimization reproduces the result of Lahini et al. almost exactly, spontaneously recovering the decoupling condition $J_{23} \approx 0$ that Ralph et al. impose by construction, and saturating the theoretical bound with $\mathcal{F} = 0.9962$ at $p_{\text{succ}} \approx 1/9$. Beyond this point, the fidelity degrades smoothly and settles onto a robust floor at $\mathcal{F} = 1/\sqrt{2}$, which was shown to be a genuine consequence of the factorizing into an uncorrelated product evolution, a limit independently confirmed by the analogous plateau found in the results of Scala et al. under a different fidelity convention. The qualitatively different, monotonically decreasing tail reported by Lahini et al., however, could not be reconciled with our results and remains an open question.

Switching on the nonlinearity reshapes this picture substantially, though not uniformly. Near the linear-optical bound, $\Gamma = 0$ remains the best-performing configuration found in this work, and moderate nonlinearities offer no advantage there. The benefit of Γ instead emerges at high success probability, where linear optics is fundamentally incapable of sustaining high fidelity: at $\Gamma = 0.5$ and $\Gamma = 0.75$, this is achieved only partially, through an unresolved competition between multiple, qualitatively distinct decoupling families, producing markedly irregular Pareto fronts. At $\Gamma = 1$, this competition resolves in favour of a single robust family, sustaining $\mathcal{F} > 0.997$ for $p_{\text{succ}} \gtrsim 0.83$ and approaching fully deterministic operation as $p_{\text{succ}} \rightarrow 1$. Across all values of Γ , the mechanism enabling fidelities beyond the linear floor was traced, at the level of individual optimized parameters, to qualitatively different hopping topologies: from the near-decoupling of a single waveguide pair, to a three-site "engine" conditionally accessed by the control-qubit photon, to the essentially full recombination onto the logical subspace found at $\Gamma = 1$.

Finally, an order-of-magnitude comparison with the physical estimates of Scala et al. showed that the normalized nonlinearity strengths explored in this chapter, Γ/J_{max} , are broadly compatible with values considered experimentally achievable in exciton-polariton waveguide platforms, lending some physical plausibility to the results presented, although a definitive assessment would require committing to a specific material platform. Taken together, these results indicate that a genuinely deterministic CNOT gate is, in principle, achievable within the same minimal six-waveguide geometry originally proposed for the probabilistic linear-optical case, provided a sufficiently strong and continuously distributed self-Kerr nonlinearity is available – without resorting to the multi-block concatenated architecture required by previous nonlinear proposals.

Conclusions

This thesis started from asking ourselves a relatively simple question: can the architectural simplicity of the quantum-walk CNOT gate proposed by Y. Lahini et al. a few years ago [15] be combined with a distributed self-Kerr non-linearity, in order to overcome the $1/9$ success-probability bound that is well known from the linear optical quantum computing paradigm, without giving up the minimal six-waveguide geometry? In fact, Chapter 4 answered this question in a systematic way, and the final answer is: yes, it does.

To briefly summarize the main outcomes of this work, we started from the analysis of the purely linear quantum walk case, i.e., by imposing $\Gamma = 0$ (Γ representing the energy scale associated with the two-particle nonlinearity within the propagating system, given in dimensionless units). Here, the optimization reproduces the result of Lahini et al. almost exactly: the algorithm spontaneously finds the decoupling condition $J_{23} \approx 0$ that Ralph et al. [10] impose by construction, and it saturates the theoretical bound with fidelity $F = 0.9962$ at $p_{\text{succ}} \approx 1/9$. At difference with the result previously found by Lahini et al., in our numerical optimization the fidelity does not collapse when moving away from this point. Instead, it settles smoothly onto a robust floor at $F = 1/\sqrt{2}$. We showed that this floor is not a numerical artefact, but rather a real consequence of the system factorizing into an uncorrelated product evolution on the two qubits. This same plateau also appears, independently, in the results of Scala et al. [17], under a different fidelity convention, which supports our interpretation. One thing we could not explain is the different, monotonically decreasing tail reported by Lahini et al.: this remains an open point for future work.

Switching on the non-linearity (i.e., setting $\Gamma \neq 0$ since the beginning of the numerical optimization procedure) radically changes this picture, albeit smoothly and non-uniformly. Close to the linear-optical bound, $\Gamma = 0$ is still the best choice we found: a small or moderate non-linearity gives no advantage there. The real benefit of Γ appears at high success probability, exactly where linear optics cannot keep the fidelity high. At $\Gamma = 0.5$ and $\Gamma = 0.75$ (for the meaning of these dimensionless values, we remind the reader to the presentation of results in Chapter 4), we found that this improvement is only partial, and comes from a competition between multiple, qualitatively different families of solutions, which makes the Pareto front irregular. At $\Gamma = 1$, this competition resolves: a single, robust family of solutions dominates the whole front, keeping $F > 0.997$ for $p_{\text{succ}} \gtrsim 0.83$, and approaching a fully deterministic gate as $p_{\text{succ}} \rightarrow 1$. Looking at the individual optimized parameters, we could trace this improvement to different hopping topologies: from the near-decoupling of a single pair of waveguides, to a three-site "engine" conditionally used by the control-qubit photon, to an almost complete recombination onto the logical subspace at $\Gamma = 1$. This is considered the main

result of the thesis.

Furthermore, we also asked whether these results are physically reasonable, not just mathematically optimal. Using an order-of-magnitude estimate based on the physical scale introduced in Sec. 4.6, we compared our dimensionless non-linearity values, i.e., normalized to the largest hopping rate found by the optimizer during each different iteration (i.e., $\Gamma/J_{\max}$ ranging from 2.2×10^{-3} to 0.125) with the values that Scala et al. report as realistic to optimistic for exciton-polariton waveguide platforms ($U/J_{\max} \approx 0.01\text{--}0.5$). All our results fall inside this physically motivated range. This agreement was not built into the optimization in any way, so it gives some real physical credibility to our non-linear results, even if a definitive answer would require choosing one specific material platform.

Taken altogether, these results show that a genuinely deterministic CNOT gate is achievable within the same minimal six-waveguide geometry originally used for the probabilistic, linear-optical case, as long as a sufficiently strong and uniformly distributed self-Kerr non-linearity is available. This is a simpler route to determinism than the multi-block architecture required by previous non-linear proposals [17]. The results of Chapter 4 open more questions than they close, and several natural directions follow directly from this work.

The most direct one is experimental. An exciton-polariton implementation of a non-linear photonic CNOT gate is currently being explored by other groups working on this platform [17, 69], and the optimized parameters found in this thesis could serve as a concrete target for such an effort. To make this connection realistic, the model used here would need to be extended: dissipation should be added through Lindblad master equations, to model photon loss in a realistic semiconductor device, and full electromagnetic-field propagation should be simulated with a Maxwell solver, to include dispersion and realistic waveguide geometry. The optimization itself could also be extended to richer non-linearity structures, such as site-dependent self-Kerr terms or a combination of self-Kerr and cross-Kerr interactions, and to other gate geometries and other two- or three-qubit gates (for example CZ or Toffoli gates). More broadly, this optimization-based approach connects to the growing field of quantum optical neural networks, where trainable nets of linear optics and single-site non-linearities are used to design a wide range of quantum operations [18, 19]; extending the present six-waveguide geometry into such a trainable network could offer a natural path toward the higher-dimensional, multi-gate architectures mentioned above.

A second, more theoretical direction is to understand why the optimized system works, rather than only that it works. The numerical optimizer finds high-fidelity configurations, but treats the system as a black box. Some features of the optimized Hamiltonians found in Chapter 4 suggest that something structured is happening: the two-photon dynamics is regular and close to periodic, the ancilla waveguides are effectively excluded at high Γ , and each logical state is built from only a few dominant eigenstates instead of being spread over the whole system. This raises a natural question: does the probabilistic-to-deterministic transition, driven by Γ , correspond to a genuine spectral phase transition of the optimized Hamiltonian? Two complementary pictures could be used to address this. The first is many-body localization, where the transition would separate a thermalizing phase, in which the photons explore the full system including the ancillas, from a localized phase, in which the dynamics stays confined to the logical subspace. The second is the theory of Hopfield networks, where the inhomogeneous

couplings J_{ij} found by the optimizer could be seen as storing the gate as a stable dynamical attractor. If this picture holds, it would also suggest a natural generalization: instead of a one-dimensional chain of six waveguides, one could design a disordered, higher-dimensional network of coupled waveguides, engineered to store several quantum gates at once as distinct attractors, moving toward a programmable non-linear photonic quantum processor.

Appendix A

Simulation code

This appendix reports the Python implementation of the numerical simulations discussed in Chapter 4. The code excerpt reported below represents only a portion of the full implementation; auxiliary parts that are not essential for understanding the algorithm, but are merely required for its execution, have been omitted for clarity. The first script performs a single optimization run, scanning the full set of target success probabilities and producing one realization of the fidelity–success-probability Pareto front for a given value of Γ . The second script aggregates the results of multiple independent runs of the first script, selecting, for each target probability, the configuration with the highest fidelity, in order to obtain the consolidated Pareto front and the associated physical plots discussed in the main text. Both of the codes take the case of $\Gamma = 1$ as the paradigmatic example.

A.1 Single-run optimization script

Listing A.1: Optimization script (single run)

```
1 #=====
2 #Space settings
3 #=====
4
5 #Working space dimension (in our case, for two relevant qubits, it is
6   4)
7 NSpace=4
8
9 # Time and nonlinearity
10 T = 1
11 Gamma= 1
12
13 # Define a, adaga and Id as dense 3x3 matrices initially
14 a_dense=np.array([[0, np.sqrt(1), 0],
15                  [0, 0, np.sqrt(2)],
16                  [0, 0, 0]])
17 adaga_dense=np.array([[0, 0, 0],
18                      [np.sqrt(1), 0, 0],
19                      [0, np.sqrt(2), 0]])
20 Id_dense=np.array([[1,0,0],
21                   [0,1,0],
```

```

21         [0,0,1]])
22
23 # Convert the matrices to sparse format (CSR - Compressed Sparse Row)
24 a = sparse.csr_matrix(a_dense)
25 adaga = sparse.csr_matrix(adaga_dense)
26 Id = sparse.csr_matrix(Id_dense)
27 n=adaga@a          #automatically a sparse matrix
28
29 #Now extend the single a and adaga to the remaining subspaces using
    sparse.kron
30 A_1_daga= sparse.kron(adaga, sparse.kron(Id, sparse.kron(Id, sparse.
    kron(Id, sparse.kron(Id, Id))))))
31 A_2_daga= sparse.kron(Id, sparse.kron(adaga, sparse.kron(Id, sparse.
    kron(Id, sparse.kron(Id, Id))))))
32 A_3_daga= sparse.kron(Id, sparse.kron(Id, sparse.kron(adaga, sparse.
    kron(Id, sparse.kron(Id, Id))))))
33 A_4_daga= sparse.kron(Id, sparse.kron(Id, sparse.kron(Id, sparse.kron(
    adaga, sparse.kron(Id, Id))))))
34 A_5_daga= sparse.kron(Id, sparse.kron(Id, sparse.kron(Id, sparse.kron(
    Id, sparse.kron(adaga, Id))))))
35 A_6_daga= sparse.kron(Id, sparse.kron(Id, sparse.kron(Id, sparse.kron(
    Id, sparse.kron(Id, adaga))))))
36
37 A_1= sparse.kron(a, sparse.kron(Id, sparse.kron(Id, sparse.kron(Id,
    sparse.kron(Id, Id))))))
38 A_2= sparse.kron(Id, sparse.kron(a, sparse.kron(Id, sparse.kron(Id,
    sparse.kron(Id, Id))))))
39 A_3= sparse.kron(Id, sparse.kron(Id, sparse.kron(a, sparse.kron(Id,
    sparse.kron(Id, Id))))))
40 A_4= sparse.kron(Id, sparse.kron(Id, sparse.kron(Id, sparse.kron(a,
    sparse.kron(Id, Id))))))
41 A_5= sparse.kron(Id, sparse.kron(Id, sparse.kron(Id, sparse.kron(Id,
    sparse.kron(a, Id))))))
42 A_6= sparse.kron(Id, sparse.kron(Id, sparse.kron(Id, sparse.kron(Id,
    sparse.kron(Id, a))))))
43
44 N_1 = A_1_daga @ A_1
45 N_2 = A_2_daga @ A_2
46 N_3 = A_3_daga @ A_3
47 N_4 = A_4_daga @ A_4
48 N_5 = A_5_daga @ A_5
49 N_6 = A_6_daga @ A_6
50
51 # List of number operators
52 N_list = [N_1, N_2, N_3, N_4, N_5, N_6]
53
54 # List of annihilation operators
55 A_list = [A_1, A_2, A_3, A_4, A_5, A_6]
56
57 # List of creation operators
58 A_dag_list = [A_1_daga, A_2_daga, A_3_daga, A_4_daga, A_5_daga,
    A_6_daga]
59
60 I = sparse.eye(N_list[0].shape[0], format='csr')
61
62 ketVuoto=np.zeros(729)          #vector of one one followed by 728 zeros

```



```

63 ketVuoto[0]=1
64
65 #logical states
66 ket00_dense=(A_2_daga@(A_4_daga@ketVuoto))
67 ket01_dense=(A_2_daga@(A_5_daga@ketVuoto))
68 ket10_dense=(A_3_daga@(A_4_daga@ketVuoto))
69 ket11_dense=(A_3_daga@(A_5_daga@ketVuoto))
70
71 hopping_ops = [A_dag_list[i] @ A_list[i+1] + A_dag_list[i+1] @ A_list[
    i] for i in range(5)]
72 basis_matrix = np.column_stack([ket00_dense, ket01_dense, ket10_dense,
    ket11_dense])
73 basis_matrix_conj = basis_matrix.conj().T
74 CNOT = np.array([[1,0,0,0],[0,1,0,0],[0,0,0,1],[0,0,1,0]], dtype=
    complex)
75 CNOT_daga = CNOT.conj().T
76
77 #=====
78 # Optimizer configuration
79 #=====
80 n_parametri = 10
81 lowEJ=-8          #minimum bound value
82 uppEJ =8          #maximum bound value
83 w=20
84
85 # =====
86 # UTILITY FUNCTION: builds H and computes the 4x4 gate
87 # =====
88 def compute_gate(parametri):
89     E_i = np.zeros(6)
90     J_ij = np.zeros(5)
91     E_i[1:6] = parametri[0:5]
92     E_i[0] = 0.0
93     J_ij[0:5] = parametri[5:10]
94
95 #Hamiltonian
96 H_full = (sum(E_i[i] * N_list[i] for i in range(6)) +
97           sum(J_ij[i] * hopping_ops[i] for i in range(5))) + Gamma*
98           sum(N_list[i]@(N_list[i]-I) for i in range(6))
99
100 #Projected matrix
101 evolved_matrix = expm_multiply(-1j * np.pi * T * H_full,
    basis_matrix)
102 G4x4 = basis_matrix_conj @ evolved_matrix
103 return H_full, G4x4
104
105 # =====
106 #COST FUNCTION DEFINITION
107 # =====
108 def CostF(parametri):
109
110     #Energy of the 6 sites
111     E_i = np.zeros(6)
112     #Tunneling rate between the 6 sites (j=i+1, i.e.: ij
    =[01,12,23,34,45])

```

```

113 J_ij=np.zeros(5)
114
115 #Unpack parameters from the input array
116 E_i[1], E_i[2], E_i[3], E_i[4], E_i[5], J_ij[0], J_ij[1], J_ij[2],
    J_ij[3], J_ij[4]= parametri
117 E_i[0]= 0.00
118
119 #HAMILTONIAN and gate (delegated to compute_gate)
120 _, G4x4 = compute_gate(parametri)
121
122 #Normalize and then remove the phase
123 GdagaG = G4x4.conj().T @ G4x4
124 eigvals= np.linalg.eigvalsh(GdagaG)
125 succes_prob_local = np.mean(eigvals)
126
127 G4x4_normalized = G4x4 / np.linalg.norm(G4x4, 'fro')
128
129 Fidelity=np.abs(np.trace(CNOT.conj().T @ G4x4_normalized)/ np.linalg
    .norm(CNOT, 'fro'))
130
131 return Fidelity, succes_prob_local, G4x4, G4x4_normalized
132 # =====
133 #Cost Function
134 # =====
135 def nlopt_objective(x, grad):
136     Fidelity, succes_prob_local, G4x4 , G4x4_normalized = CostF(x)
137     objective_value = float(1-Fidelity**2)
138     penalty = w*(succes_prob_local - probab_scelta)**2
139
140     return objective_value + penalty
141
142 # =====
143 # NLopt OPTIMIZER: global MLSL + local BOBYQA
144 # =====
145 nomeopt = "MLSL"
146
147 # Local optimizer used INTERNALLY by MLSL during the global search
148 opt_local_for_msls = nlopt.opt(nlopt.LN_BOBYQA, n_parametri)
149 opt_local_for_msls.set_lower_bounds([lowEJ] * n_parametri)
150 opt_local_for_msls.set_upper_bounds([uppEJ] * n_parametri)
151 opt_local_for_msls.set_xtol_rel(1e-3)
152 opt_local_for_msls.set_maxeval(20)
153
154 # Global MLSL optimizer
155 opt_global = nlopt.opt(nlopt.G_MLSL_LDS, n_parametri)
156 opt_global.set_lower_bounds([lowEJ] * n_parametri)
157 opt_global.set_upper_bounds([uppEJ] * n_parametri)
158 opt_global.set_min_objective(nlopt_objective)
159 opt_global.set_local_optimizer(opt_local_for_msls)
160 opt_global.set_maxeval(18000)
161 opt_global.set_xtol_rel(1e-2)
162
163 # Local optimizer for the FINAL refinement on the best global point
164 opt_refine = nlopt.opt(nlopt.LN_BOBYQA, n_parametri)
165 opt_refine.set_lower_bounds([lowEJ] * n_parametri)
166 opt_refine.set_upper_bounds([uppEJ] * n_parametri)

```

```

167 opt_refine.set_min_objective(nlopt_objective)
168 opt_refine.set_xtol_rel(1e-6)
169 opt_refine.set_maxeval(1500)
170 opt = opt_global
171
172 # =====
173 # TARGET PROBABILITY
174 # =====
175
176 intervalli = 80
177 valoridadare_prob_scelta = np.array(np.linspace(0.02, 0.99, intervalli
178 ))
179 print("Genero prob_scelta_values:")
180 print(valoridadare_prob_scelta)
181
182 #=====
183 #START OPTIMIZATION
184 #=====
185 probabilita_successo = np.zeros(intervalli)
186 valori_fidelity = np.zeros(intervalli)
187 valori_costo = np.zeros(intervalli)
188 valori_parametri = np.zeros((intervalli, n_parametri))
189
190 # Define the bounds for the random generation of x0
191 lower_bounds_x0 = [lowEJ] * (n_parametri)
192 upper_bounds_x0 = [uppEJ] * (n_parametri)
193
194 rng = np.random.default_rng()
195 random_seed = rng.integers(0, 50)
196 np.random.seed(random_seed)
197 for i in range(intervalli):
198     prob_scelta = valoridadare_prob_scelta[i]
199     print(f"\nOttimizzatore {nomeopt}, iterazione {i+1} per prob target
200           = {prob_scelta}...")
201
202     # Random starting point
203     x0 = np.random.uniform(lower_bounds_x0, upper_bounds_x0, n_parametri
204 )
205     try:
206         x_opt = opt_global.optimize(x0)
207         x_final = opt_refine.optimize(x_opt)
208         # Retrieve the final data
209         Fidelity, succes_prob_local, G4x4, G4x4_normalized = CostF(x_final
210 )
211
212         costo_minimo = opt_refine.last_optimum_value()
213         valori_parametri[i, :] = x_final
214         valori_costo[i] = costo_minimo
215         valori_fidelity[i] = Fidelity
216         probabilita_successo[i] = succes_prob_local
217
218     except nlopt.RoundoffLimited:
219         print(" Attenzione: limite di precisione raggiunto.")
220     except Exception as e:

```

```

219     print(f" Errore: {e}")
220
221 print("fine iterazioni")

```

A.2 Multi-run aggregation script

Listing A.2: Aggregation script for multiple runs (mesh up)

```

1  # =====
2  # MASHUP: best fidelity for each prob_scelta
3  # =====
4
5  # All unique prob_scelta values present in the loaded files
6  tutti_prob_scelta = np.unique(np.round(
7      np.concatenate([d["prob_scelta"] for d in tutti_i_dati]), 10))
8
9  print(f"\nUnique prob_scelta values: {len(tutti_prob_scelta)}")
10 print(f"Range: [{tutti_prob_scelta.min():.4f}, {tutti_prob_scelta.max()
11     ():.4f}]")
12
13 n_prob = len(tutti_prob_scelta)
14 n_parametri = tutti_i_dati[0]["valori_parametri"].shape[1]
15
16 best_prob_scelta = np.zeros(n_prob)
17 best_fidelity = np.zeros(n_prob)
18 best_prob_successo = np.zeros(n_prob)
19 best_costo = np.zeros(n_prob)
20 best_parametri = np.zeros((n_prob, n_parametri))
21 best_file_origine = []
22
23 print("\nSelecting best fidelity for each prob_scelta...")
24
25 for idx, p_target in enumerate(tutti_prob_scelta):
26     fidelity_migliore = -1.0
27     best_entry = None
28     file_migliore = "N/A"
29
30     for dati in tutti_i_dati:
31         mask = np.abs(dati["prob_scelta"] - p_target) <=
32             TOLLERANZA_PROB
33         if not np.any(mask):
34             continue
35
36         fidelity_candidate = dati["valori_fidelity"][mask]
37         idx_locale = np.argmax(fidelity_candidate)
38         idx_globale = np.where(mask)[0][idx_locale]
39
40         if fidelity_candidate[idx_locale] > fidelity_migliore:
41             fidelity_migliore = fidelity_candidate[idx_locale]
42             best_entry = {
43                 "fidelity": dati["valori_fidelity"][idx_globale],
44                 "prob_successo": dati["probabilita_successo"][
45                     idx_globale],
46                 "costo": dati["valori_costo"][idx_globale],

```

```

44         "parametri":      dati["valori_parametri"][idx_globale
45         ],
46     }
47     file_migliore = dati["nome_file"]      # correct traceability
48
49     if best_entry is None:
50         print(f" [WARN] No entry for prob_scelta={p_target:.6f}")
51         best_prob_scelta[idx] = p_target
52         best_fidelity[idx] = np.nan
53         best_prob_successo[idx] = np.nan
54         best_costo[idx] = np.nan
55         best_parametri[idx] = np.full(n_parametri, np.nan)
56     else:
57         best_prob_scelta[idx] = p_target
58         best_fidelity[idx] = best_entry["fidelity"]
59         best_prob_successo[idx] = best_entry["prob_successo"]
60         best_costo[idx] = best_entry["costo"]
61         best_parametri[idx] = best_entry["parametri"]
62     best_file_origine.append(file_migliore)
63
64 print(f"Selection completed for {n_prob} prob_scelta values.")
65
66 #=====
67 #PARETO FRONT PLOT
68 #=====
69
70 plt.figure(figsize=(8, 4))
71 plt.plot(best_prob_successo[valid], best_fidelity[valid], 'o-',
72         markersize=5, linewidth=1.2, label="Best fidelity per
73         prob_scelta")
74 plt.axvline(x=1/9, color='red', linestyle='--', linewidth=1, label=f"
75         Limite teorico 1/9 =1/9:.4f)")
76 plt.xlabel("Success Probability")
77 plt.ylabel("Fidelity")
78 plt.title(f"Pareto front: fidelity vs success probability {name_graph}
79         ")
80 plt.legend()
81 plt.grid(True)
82 plt.tight_layout()
83 plt.savefig(f"best_fidelity_vs_prob_{name_graph}.png", dpi=150)
84 plt.show()
85
86 #=====
87 # PHYSICAL PLOTS FOR THE SELECTED ITERATION
88 #=====
89
90 idx_plot = idx_max if ITERAZIONE_SCELTA is None else int(
91     ITERAZIONE_SCELTA)
92
93 print(f"\n--- Physical analysis: iteration {idx_plot} ---")
94 print(f"prob_scelta:      {best_prob_scelta[idx_plot]:.6f}")
95 print(f"Fidelity:         {best_fidelity[idx_plot]:.6f}")
96 print(f"Prob_successo:      {best_prob_successo[idx_plot]:.6f}")
97 print(f"Costo:              {best_costo[idx_plot]:.6f}")
98 print(f"Source file:        {Path(best_file_origine[idx_plot]).name}")

```

```

94 print(f"Parameters:      {np.array2string(best_parametri[idx_plot],
95     separator=', ', max_line_width=np.inf)}")
96
97 selected_params = best_parametri[idx_plot]
98 H_full, G4x4 = compute_gate(selected_params)
99
100 GdagaG          = G4x4.conj().T @ G4x4
101 eigvals         = np.linalg.eigvalsh(GdagaG)
102 succes_prob_local = np.mean(eigvals)
103 G4x4_normalized  = G4x4 / np.linalg.norm(G4x4, 'fro')
104 phi_global       = np.angle(G4x4_normalized)
105 G4x4_phase_corrected = G4x4_normalized * np.exp(-1j * phi_global)
106
107 print("\nG4x4:")
108 print(np.round(G4x4, 4))
109 print("\nNormalized G4x4:")
110 print(np.round(G4x4_normalized, 4))
111 print("\nNormalized G4x4, phase removed:")
112 print(np.round(G4x4_phase_corrected, 4))
113
114 # =====
115 # 3D plot -- Real and Imaginary part of G4x4_phase_corrected
116 # =====
117 real_G4x4      = np.real(G4x4_phase_corrected)
118 imag_G4x4      = np.imag(G4x4_phase_corrected)
119 state_labels    = ['|00>', '|01>', '|10>', '|11>']
120
121 fig, (ax1, ax2) = plt.subplots(1, 2, figsize=(12, 5), subplot_kw={'
122     projection': '3d'})
123
124 x_pos = np.arange(4)
125 y_pos = np.arange(4)
126 X_pos, Y_pos = np.meshgrid(x_pos, y_pos)
127 X_flat = X_pos.flatten()
128 Y_flat = Y_pos.flatten()
129 Z_flat = np.zeros_like(X_flat)
130 dx = dy = 0.7
131
132 ax1.bar3d(X_flat, Y_flat, Z_flat, dx, dy, real_G4x4.flatten(), cmap='
133     viridis', zsort='average')
134 ax1.set_xticks(x_pos); ax1.set_xticklabels(state_labels, rotation
135     =-15, ha='left')
136 ax1.set_yticks(y_pos); ax1.set_yticklabels(state_labels, rotation=15,
137     ha='left')
138 ax1.set_xlabel('Input'); ax1.set_ylabel('Output')
139 ax1.set_zlabel('Re(G)'); ax1.set_title('Real Part')
140 ax1.set_zlim(-0.1, 1.1)
141
142 ax2.bar3d(X_flat, Y_flat, Z_flat, dx, dy, imag_G4x4.flatten(), cmap='
143     plasma', zsort='average')
144 ax2.set_xticks(x_pos); ax2.set_xticklabels(state_labels, rotation
145     =-15, ha='left')
146 ax2.set_yticks(y_pos); ax2.set_yticklabels(state_labels, rotation=15,
147     ha='left')
148 ax2.set_xlabel('Input'); ax2.set_ylabel('Output')
149 ax2.set_zlabel('Im(G)'); ax2.set_title('Imaginary Part')

```

```

142 ax2.set_zlim(-0.08, 0.08)
143
144 fig.suptitle(f"G4x4 - iterazione {idx_plot}_{name_graph}(F={
    best_fidelity[idx_plot]:.4f})", fontsize=13)
145 plt.tight_layout()
146 plt.savefig(f"G4x4_3d_plot_iter_{idx_plot}_{name_graph}.png", dpi=150)
147 plt.show()
148 files.download(f"G4x4_3d_plot_iter_{idx_plot}_{name_graph}.png")
149
150 # =====
151 # GRADIENT PLOTS -- correct physical normalization
152 # =====
153 initial_states_mapping = {
154     "00": np.array(ket00).flatten(),
155     "01": np.array(ket01).flatten(),
156     "10": np.array(ket10).flatten(),
157     "11": np.array(ket11).flatten(),
158 }
159
160 times = np.linspace(0, T, 100)
161
162 fig_grad, axes_grad = plt.subplots(1, 4, figsize=(13, 4))
163 fig_grad.suptitle('Density Maps for Different Initial States',
    fontsize=16)
164
165 for idx, (label, psi0) in enumerate(initial_states_mapping.items()):
166     ax = axes_grad[idx]
167     density = np.zeros((n_sites, len(times)))
168
169     for i, t in enumerate(times):
170         psi_t = scipy.sparse.linalg.expm_multiply(-1j * np.pi * H_full
            * t, psi0)
171         psi_t = psi_t / np.linalg.norm(psi_t) # physical
            normalization
172
173         for site_idx in range(n_sites):
174             density[site_idx, i] = np.real(
175                 psi_t.conj() @ (N_list[site_idx] @ psi_t)
176             )
177
178         im = ax.imshow(density.T, aspect='auto', origin='lower', cmap='
            inferno',
179             extent=[0.5, 6.5, times[0], times[-1]],
180             vmin=0, vmax=2)
181         ax.set_title(f"Input State {label}")
182         ax.set_xlabel('Site')
183         ax.set_ylabel('Time (T)')
184         fig_grad.colorbar(im, ax=ax)
185
186 plt.tight_layout(rect=[0, 0.03, 1, 0.95])
187 plt.savefig(f"gradient_plot_iter_{idx_plot}_{name_graph}.png", dpi
    =150)
188 plt.show()
189 files.download(f"gradient_plot_iter_{idx_plot}_{name_graph}.png")
190
191 # =====

```



```

192 # BAR PLOT at T=1 -- correct physical normalization
193 # =====
194 fig_bars, axes_bars = plt.subplots(1, 4, figsize=(16, 5), sharey=True)
195 fig_bars.suptitle(f'Density at T=1 for Each Initial State_{name_graph}'
196                   ', fontsize=16)
197
198 for idx, (label, psi0) in enumerate(initial_states_mapping.items()):
199     ax = axes_bars[idx]
200
201     psi_T1 = expm_multiply(-1j * np.pi * H_full * T, psi0)
202     psi_T1 = psi_T1 / np.linalg.norm(psi_T1) # physical normalization
203
204     density_at_T1 = np.array([
205         np.real(psi_T1.conj() @ (N_list[site_idx] @ psi_T1))
206         for site_idx in range(n_sites)
207     ])
208
209     ax.bar(np.arange(1, n_sites + 1), density_at_T1, color='skyblue')
210     ax.set_title(f"Input State {label}")
211     ax.set_xlabel('Site')
212     if idx == 0:
213         ax.set_ylabel('Density (at T=1)')
214         ax.set_xticks(np.arange(1, n_sites + 1))
215         ax.grid(axis='y', linestyle='--', alpha=0.7)
216         ax.set_ylim(0, 2)
217
218 plt.tight_layout(rect=[0, 0.03, 1, 0.95])
219 plt.savefig(f"density-T=1-barplot_iter_{idx_plot}_{name_graph}.png",
220            dpi=150)
221 plt.show()

```

Appendix B

Numerical results for small non-linearity values

B.1 Non-linear regime of $\Gamma = 0.1$ ($\Gamma/J_{\max} = 0.020$)

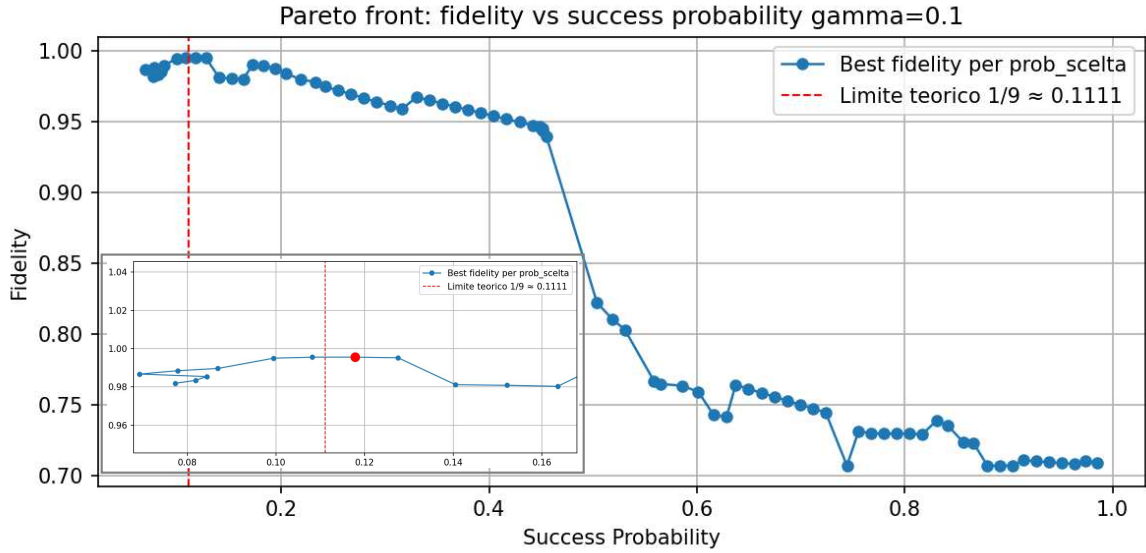


Figure B.1: Optimized Pareto front for $\Gamma = 0.1$ ($\Gamma/J_{\max} = 0.020$). The inset zooms in on the region around the highest-fidelity point, highlighted in red.

The highest-fidelity configuration is found for $p_{\text{target}} = 0.1182$. Expressing all parameters in units of $J_{\max}^{\text{loc}} = |J_{34}^*| \approx 4.699 \pi^{-1}$, the optimized parameter vector reads

$$\frac{\boldsymbol{\theta}^*}{J_{\max}^{\text{loc}}} = (0.791, -0.159, 0.598, -1.043, -0.485, -0.654, 2.7 \cdot 10^{-4}, -1.000, 0.828, -0.833), \quad (\text{B.1})$$

The resulting figures of merit are

$$p_{\text{succ}} = 0.1178, \quad \mathcal{F} = 0.9956, \quad (\text{B.2})$$

with a residual cost $C(\boldsymbol{\theta}^*) = 8.84 \cdot 10^{-3}$. The corresponding gate matrix, projected onto the logical subspace, is

$$G_{4 \times 4} = \begin{pmatrix} -0.034 - 0.369i & 0.023 + 0.013i & 0 & 0 \\ 0.023 + 0.013i & -0.010 - 0.309i & 0 & 0 \\ 0 & 0 & -0.014 + 0.001i & -0.022 - 0.343i \\ 0 & 0 & -0.022 - 0.343i & -0.020 + 0.008i \end{pmatrix}, \quad (\text{B.3})$$

and its normalized, phase-corrected version reads

$$\tilde{G} = \begin{pmatrix} 0.539 & 0.039 & 0 & 0 \\ 0.039 & 0.451 & 0 & 0 \\ 0 & 0 & 0.021 & 0.501 \\ 0 & 0 & 0.501 & 0.031 \end{pmatrix}, \quad (\text{B.4})$$

where entries below 10^{-3} have been set to zero.

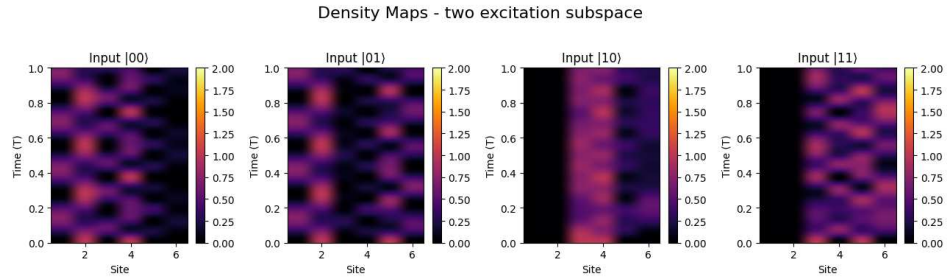


Figure B.2: Time evolution of the photon density across the six waveguides, from $T = 0$ to $T = 1$, for each computational basis input state, at the highest-fidelity point for $\Gamma = 0.1$ ($p_{\text{succ}} = 0.1178$, $\mathcal{F} = 0.9956$).

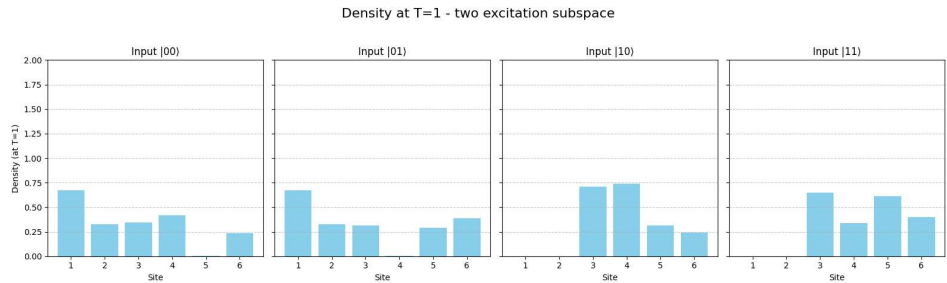


Figure B.3: Photon density at $T = 1$ across the six waveguides, for each computational basis input state, at the highest-fidelity point for $\Gamma = 0.1$ ($p_{\text{succ}} = 0.1178$, $\mathcal{F} = 0.9956$).

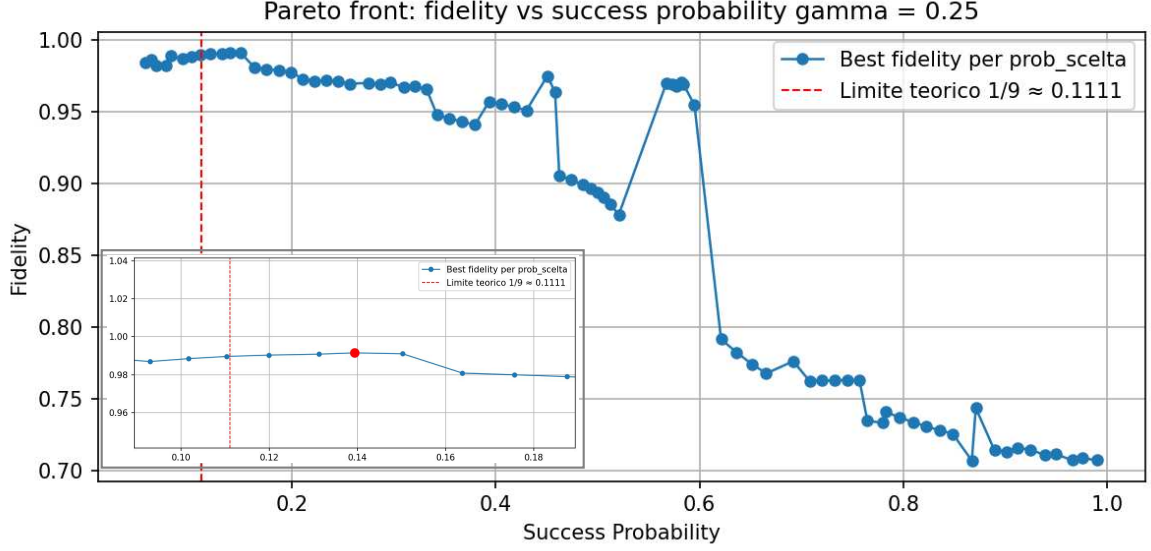


Figure B.5: Optimized Pareto front for $\Gamma = 0.25$ ($\Gamma/J_{\max} = 0.045$). The zooms in on the region around the highest-fidelity point, highlighted in red.

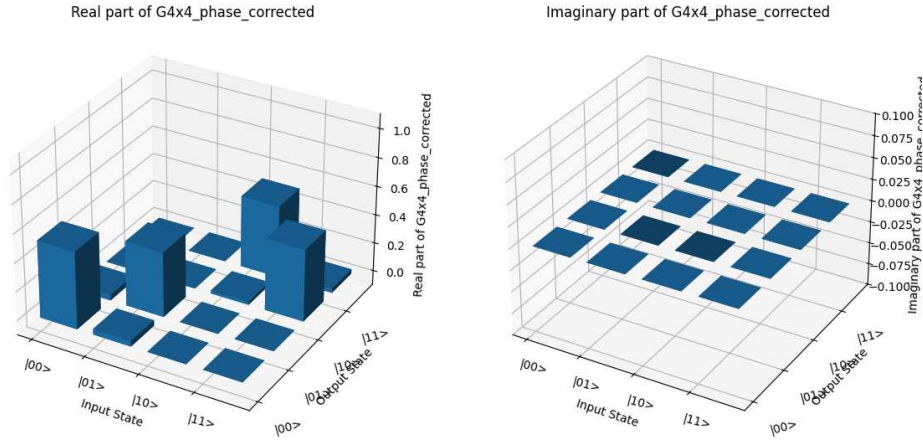


Figure B.4: Real and imaginary parts of the normalized, de-phased logical gate $\tilde{G}$ at the highest-fidelity point for $\Gamma = 0.1$ ($p_{\text{succ}} = 0.1178$, $\mathcal{F} = 0.9956$).

B.2 Non-linear regime of $\Gamma = 0.25$ ($\Gamma/J_{\max} = 0.045$)

The highest-fidelity configuration is found for $p_{\text{target}} = 0.1428$. Expressing all parameters in units of $J_{\max}^{\text{loc}} = |J_{45}^*| \approx 4.392 \pi^{-1}$, the optimized parameter vector reads

$$\frac{\boldsymbol{\theta}^*}{J_{\max}^{\text{loc}}} = (-0.455, 0.068, 0.514, -0.995, -0.626, 0.487, 4.4 \cdot 10^{-4}, -0.801, 1.000, -0.529), \quad (\text{B.5})$$

The resulting figures of merit are

$$p_{\text{succ}} = 0.1394, \quad \mathcal{F} = 0.9914, \quad (\text{B.6})$$

with a residual cost $C(\boldsymbol{\theta}^*) = 1.73 \cdot 10^{-2}$. The corresponding gate matrix, projected onto the logical subspace, is

$$G_{4 \times 4} = \begin{pmatrix} -0.383 - 0.014i & 0.050 - 0.035i & 0 & 0 \\ 0.050 - 0.035i & -0.358 - 0.051i & -0.001 & 0 \\ 0 & -0.001 & -0.004 + 0.022i & -0.367 - 0.034i \\ 0 & 0 & -0.367 - 0.034i & -0.023i \end{pmatrix}, \quad (\text{B.7})$$

and its normalized, phase-corrected version reads

$$\tilde{G} = \begin{pmatrix} 0.514 & 0.082 & 0.001 & 0 \\ 0.082 & 0.484 & 0.001 & 0.001 \\ 0.001 & 0.001 & 0.030 & 0.493 \\ 0 & 0.001 & 0.493 & 0.031 \end{pmatrix}, \quad (\text{B.8})$$

where entries below 10^{-3} have been set to zero.

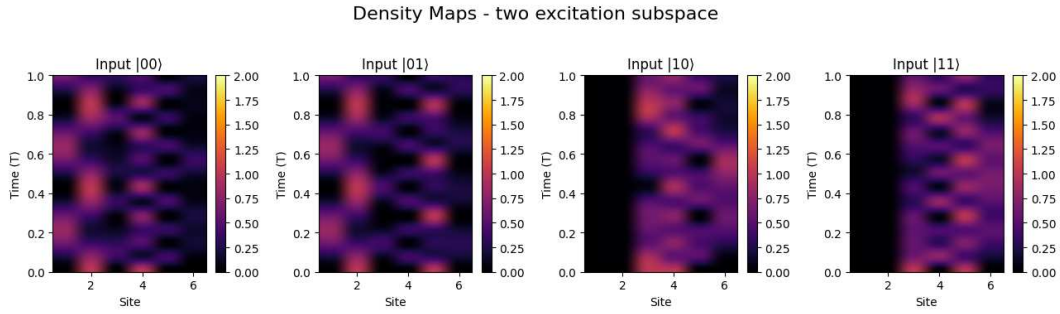


Figure B.6: Time evolution of the photon density across the six waveguides, from $T = 0$ to $T = 1$, for each computational basis input state, at the highest-fidelity point for $\Gamma = 0.25$ ($p_{\text{succ}} = 0.1394$, $\mathcal{F} = 0.9914$).

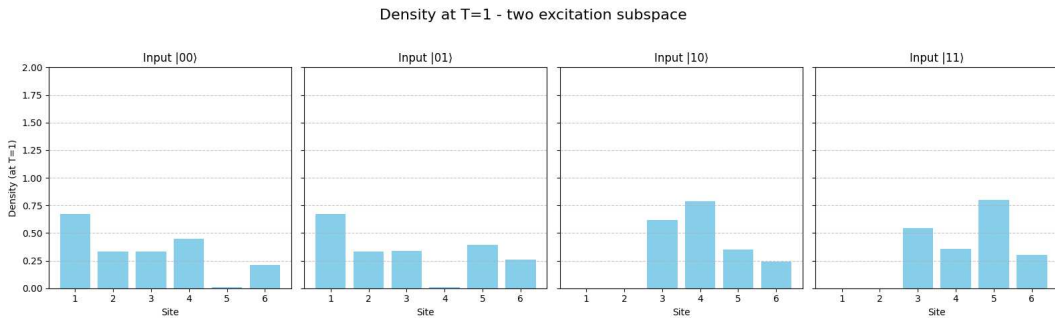


Figure B.7: Photon density at $T = 1$ across the six waveguides, for each computational basis input state, at the highest-fidelity point for $\Gamma = 0.25$ ($p_{\text{succ}} = 0.1394$, $\mathcal{F} = 0.9914$).

B.3 Non-linear regime of $\Gamma = 0.5$ ($\Gamma/J_{\text{max}} = 0.083$)

Figure B.9 shows the optimized Pareto front for $\Gamma = 0.5$. Compared to the linear case, the high-fidelity region extends visibly beyond the Ralph bound: fidelities above 0.95

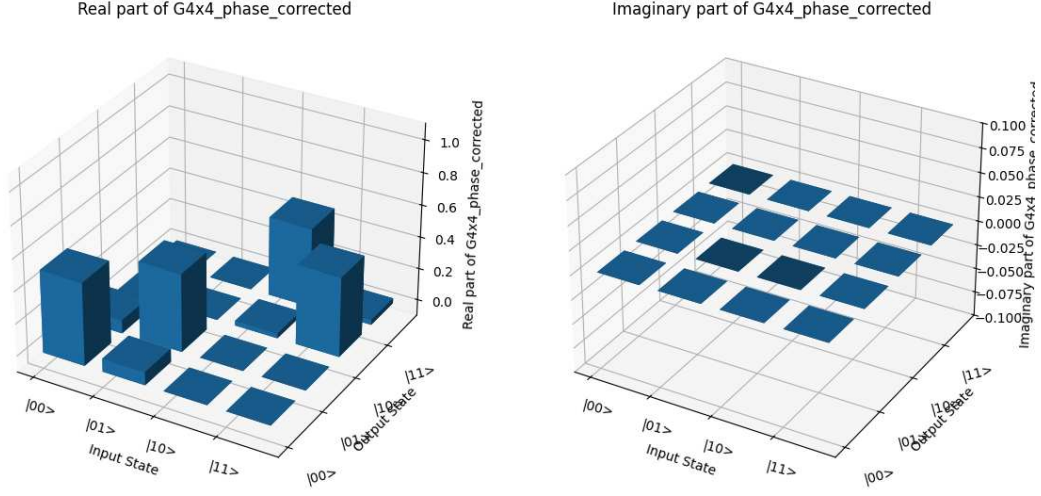


Figure B.8: Real and imaginary parts of the normalized, de-phased logical gate $\tilde{G}$ at the highest-fidelity point for $\Gamma = 0.25$ ($p_{\text{succ}} = 0.1394$, $\mathcal{F} = 0.9914$).

are sustained up to $p_{\text{succ}} \approx 0.65$, well past $p_{\text{succ}} = 1/9$. Beyond this point, however, the front becomes markedly irregular, with a sharp drop to $\mathcal{F} \approx 0.86$ around $p_{\text{succ}} \approx 0.68$, a narrow and isolated recovery to $\mathcal{F} \approx 0.975$ at $p_{\text{succ}} \approx 0.72$, and a further decline to $\mathcal{F} \approx 0.70$ as $p_{\text{succ}} \rightarrow 1$.

Highest-fidelity point The highest-fidelity configuration is found for $p_{\text{target}} = 0.1673$, close to but past the Ralph bound. Expressing all parameters in units of $J_{\text{max}}^{\text{loc}} = |J_{34}^*| \approx 5.865 \pi^{-1}$, the optimized parameter vector reads

$$\frac{\boldsymbol{\theta}^*}{J_{\text{max}}^{\text{loc}}} = (-1.019, 0.189, 0.105, -0.641, -0.749, 0.758, 3.4 \cdot 10^{-3}, 1.000, -0.716, 0.866), \quad (\text{B.9})$$

As in the linear case, $J_{23}/J_{\text{max}}^{\text{loc}} \approx 3.4 \cdot 10^{-3}$ is negligible, indicating that the spontaneous decoupling between the ancillary site $|a_1\rangle - |c_1\rangle$ block and the rest of the chain persists even in the presence of a moderate nonlinearity. The resulting figures of merit are

$$p_{\text{succ}} = 0.1657, \quad \mathcal{F} = 0.9866, \quad (\text{B.10})$$

with a residual cost $C(\boldsymbol{\theta}^*) = 0.0266$. The corresponding gate matrix, projected onto the logical subspace, is

$$G_{4 \times 4} = \begin{pmatrix} -0.307 + 0.327i & -0.007 + 0.024i & 0.002 + 0i & 0.001 - 0.002i \\ -0.007 + 0.024i & -0.310 + 0.186i & 0 + 0i & 0 + 0.001i \\ 0.002 + 0i & 0 + 0i & 0.047 + 0.024i & -0.310 + 0.256i \\ 0.001 - 0.002i & 0 + 0.001i & -0.310 + 0.256i & -0.026 - 0.055i \end{pmatrix}, \quad (\text{B.11})$$

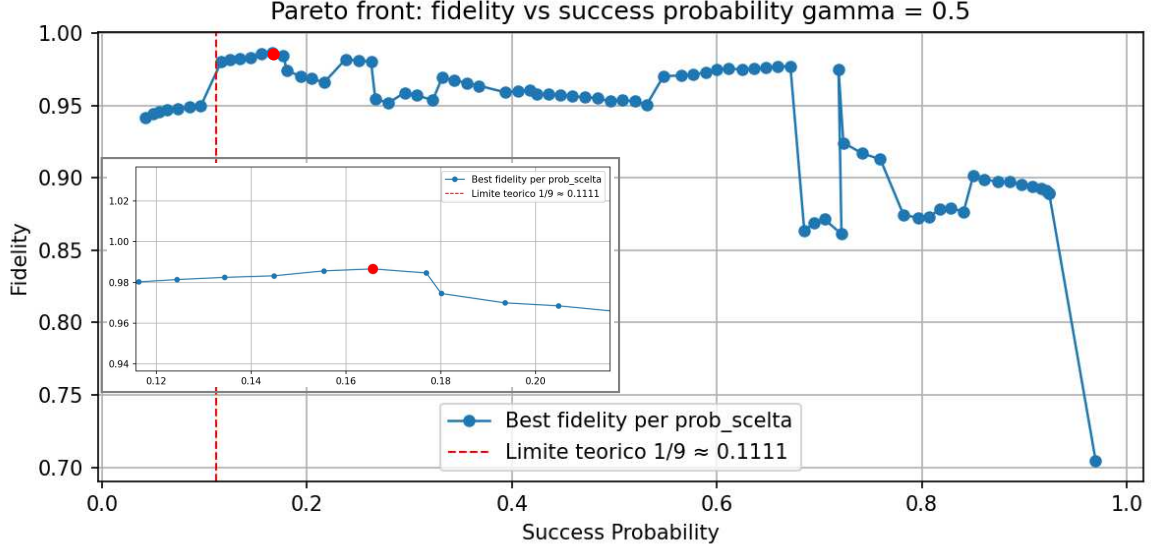


Figure B.9: Optimized Pareto front for $\Gamma = 0.5$ ($\Gamma/J_{\max} = 0.083$). The inset zooms in on the region around the highest-fidelity point, highlighted in red.

and its normalized, phase-corrected version reads

$$\tilde{G} = \begin{pmatrix} 0.551 & 0.030 & 0.003 & 0.002 \\ 0.030 & 0.445 & 0 & 0.002 \\ 0.003 & 0 & 0.065 & 0.494 \\ 0.002 & 0.002 & 0.494 & 0.075 \end{pmatrix}, \quad (\text{B.12})$$

where entries below 10^{-3} have been set to zero. The block structure inherited from the linear case is preserved, with dominant entries close to the ideal value $1/2$, consistent with the high fidelity $\mathcal{F} = 0.9866$.

Figure B.12 shows the corresponding projected onto the logical subspace, normalized, and de-phased gate $\tilde{G}$, confirming the approximate CNOT structure.

Figures B.10 and B.11 show, respectively, the time evolution of the photon density and its value at $T = 1$, for each input state of the computational basis. Unlike the linear

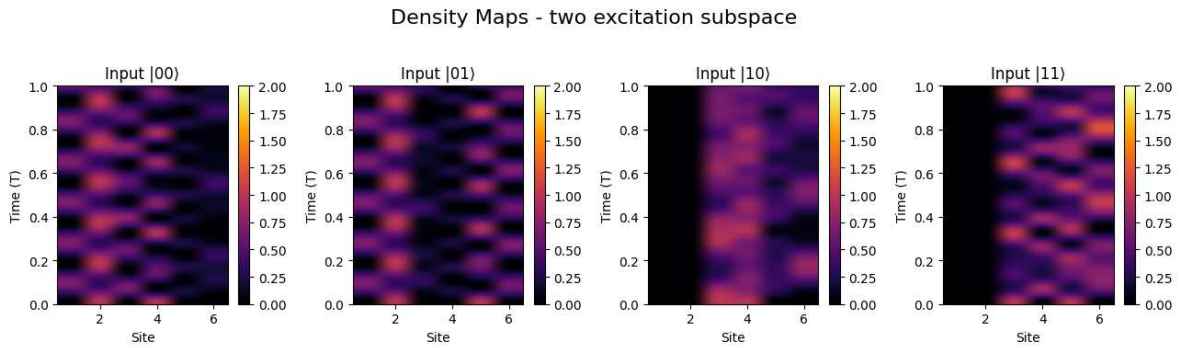


Figure B.10: Time evolution of the photon density across the six waveguides, from $T = 0$ to $T = 1$, for each computational basis input state, at the best-fidelity point for $\Gamma = 0.5$ ($p_{\text{succ}} = 0.1657$, $\mathcal{F} = 0.9866$).

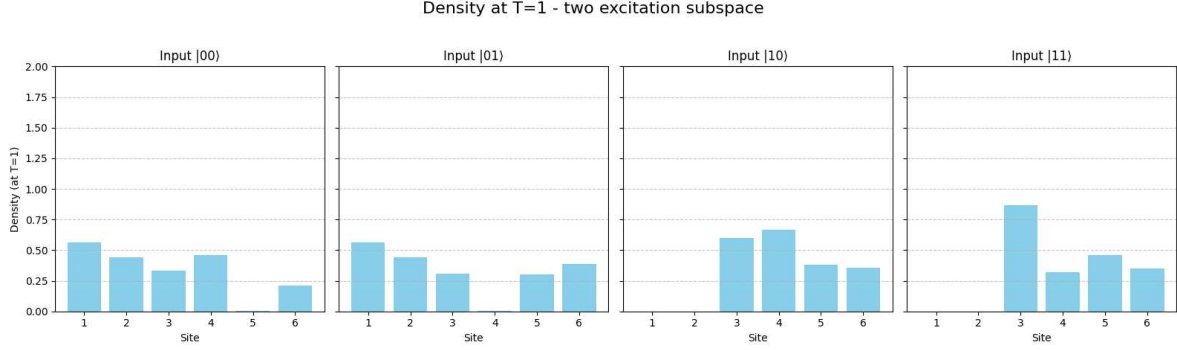


Figure B.11: Photon density at $T = 1$ across the six waveguides, for each computational basis input state, at the best-fidelity point for $\Gamma = 0.5$ ($p_{\text{succ}} = 0.1657$, $\mathcal{F} = 0.9866$).

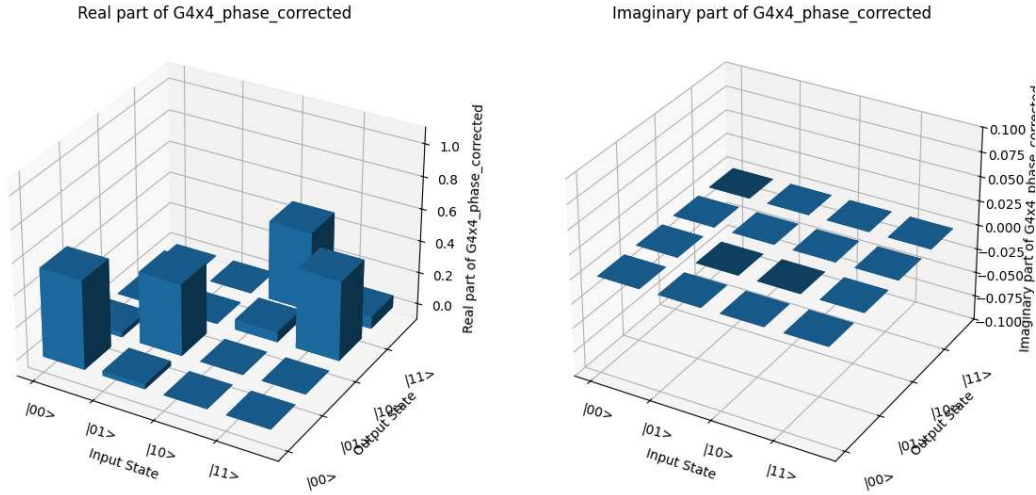


Figure B.12: Real and imaginary parts of the normalized, de-phased logical gate $\tilde{G}$ at the best-fidelity point for $\Gamma = 0.5$ ($p_{\text{succ}} = 0.1657$, $\mathcal{F} = 0.9866$).

case, the target block (sites 4–5) is no longer confined to a simple two-site oscillation: the self-Kerr term visibly reshapes the dynamics, with population spreading further into neighbouring sites during the evolution before recombining close to the ideal logical output at $T = 1$.

The isolated point of the second family To illustrate the second family of three-sites identified above, we examine the isolated recovery point at $p_{\text{target}} = 0.7322$ ($p_{\text{succ}} = 0.7188$, $\mathcal{F} = 0.9752$). Expressing the parameters in units of $J_{\text{max}}^{\text{loc}} = |J_{45}^*| \approx 5.780 \pi^{-1}$, the normalized vector reads

$$\frac{\theta^*}{J_{\text{max}}^{\text{loc}}} = (0.730, -0.120, -0.978, -1.038, -0.268, -5.2 \cdot 10^{-5}, -2.8 \cdot 10^{-7}, -0.123, 1.000, -1.5 \cdot 10^{-6}), \quad (\text{B.13})$$

Here, $J_{12}/J_{\text{max}}^{\text{loc}}$, $J_{23}/J_{\text{max}}^{\text{loc}}$, and $J_{56}/J_{\text{max}}^{\text{loc}}$ are all negligible (of order 10^{-5} – 10^{-7}), leaving only J_{34} and J_{45} as active hopping channels. The ancillary sites are fully decoupled,

and the control-qubit photon accesses the target block only when it occupies site 3. With a residual cost $C(\boldsymbol{\theta}^*) = 0.0526$, this family achieves a fidelity substantially higher than the J_{23} -only family at comparable success probability, which at this point of the front is closer to $\mathcal{F} \approx 0.86\text{--}0.87$ (see Fig. B.9).

Figures B.13 and B.14 confirm this picture: the density is confined to sites 2–5 throughout the evolution, with the $|10\rangle$ and $|11\rangle$ inputs (control qubit in $|c_2\rangle$, i.e. site 3) showing population reaching site 5, while the $|00\rangle$ and $|01\rangle$ inputs (control in $|c_1\rangle$, site 2) remain comparatively more localized near sites 2–4. The corresponding normalized, phase-corrected gate matrix,

$$\tilde{G} = \begin{pmatrix} 0.479 & 0.128 & 0 & 0 \\ 0.128 & 0.451 & 0 & 0 \\ 0 & 0 & 0.081 & 0.511 \\ 0 & 0 & 0.511 & 0.083 \end{pmatrix}, \quad (\text{B.14})$$

shown in Fig. B.15, maintains the same block-diagonal structure seen in the dominant family, with dominant entries close to the ideal reference $1/2$, consistent with the higher fidelity of this configuration. This confirms that the two competing families identified in this front, despite relying on entirely different decoupling patterns among the five hopping coefficients, converge to qualitatively similar block-structured approximations of the CNOT gate.

B.4 Non-linear regime of $\Gamma = 0.75$ ($\Gamma/J_{\max} = 0.107$)

Figure B.16 shows the optimized Pareto front for $\Gamma = 0.75$. The front is considerably more irregular than in the previous cases, repeatedly oscillating between $\mathcal{F} \approx 0.95$ and $\mathcal{F} \approx 0.98$ throughout its range, with a pronounced dip around $p_{\text{succ}} \approx 0.67$.

Highest-fidelity point The highest-fidelity configuration is found for $p_{\text{target}} = 0.1182$. Expressing all parameters in units of $J_{\max}^{\text{loc}} = |J_{45}^*| \approx 6.791 \pi^{-1}$, the optimized parameter vector reads

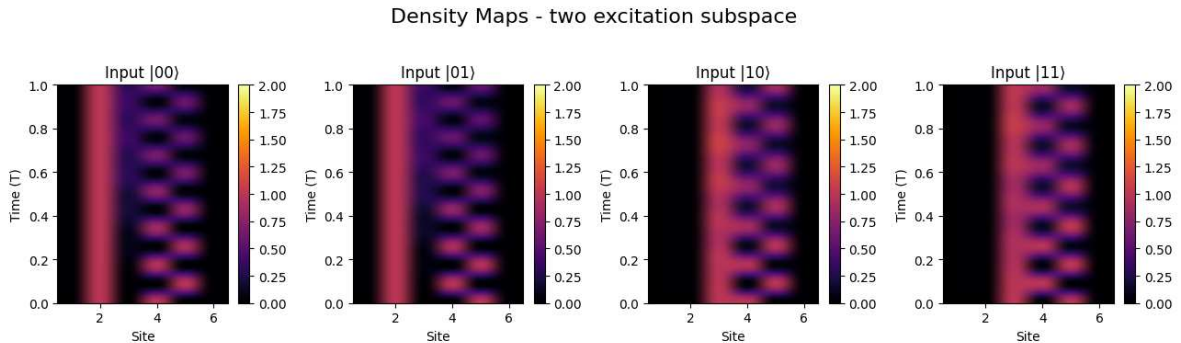


Figure B.13: Time evolution of the photon density across the six waveguides, from $T = 0$ to $T = 1$, for each computational basis input state, at the second-family point for $\Gamma = 0.5$ ($p_{\text{succ}} = 0.7188$, $\mathcal{F} = 0.9752$).

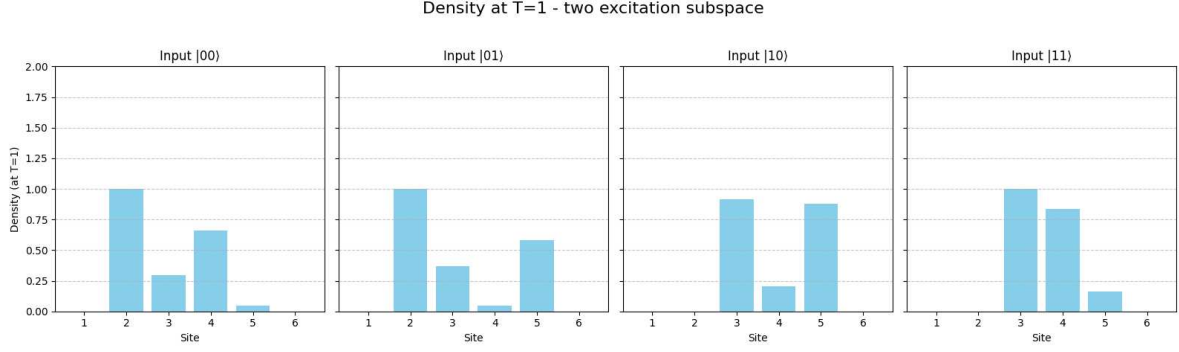


Figure B.14: Photon density at $T = 1$ across the six waveguides, for each computational basis input state, at the second-family point for $\Gamma = 0.5$ ($p_{\text{succ}} = 0.7188$, $\mathcal{F} = 0.9752$).

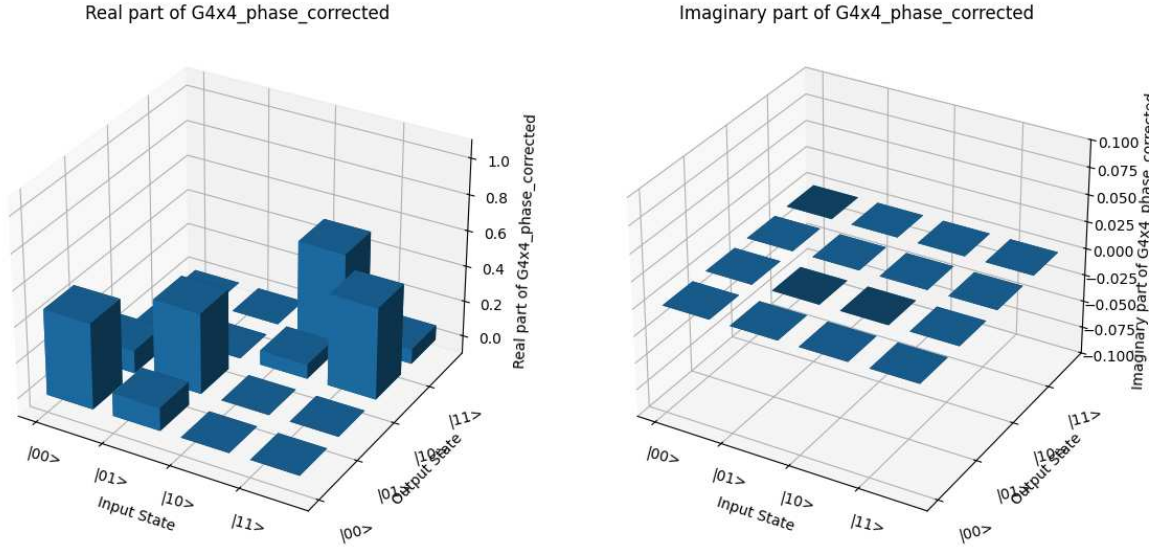


Figure B.15: Real and imaginary parts of the normalized, de-phased logical gate $\tilde{G}$ at the second-family point for $\Gamma = 0.5$ ($p_{\text{succ}} = 0.7188$, $\mathcal{F} = 0.9752$), showing the same block-diagonal structure found in the dominant family.

$$\frac{\theta^*}{J_{\text{max}}^{\text{loc}}} = (1.028, -0.531, -0.113, -1.031, -0.569, -1.000, -1.1 \cdot 10^{-3}, 0.610, 0.521, -0.541), \quad (\text{B.15})$$

As in the previous cases, $J_{23}/J_{\text{max}}^{\text{loc}} \approx -1.1 \cdot 10^{-3}$ is negligible, confirming that the decoupling between the $|a_1\rangle$ - $|c_1\rangle$ block and the rest of the chain remains a robust feature of this family. The resulting figures of merit are

$$p_{\text{succ}} = 0.1209, \quad \mathcal{F} = 0.9939, \quad (\text{B.16})$$

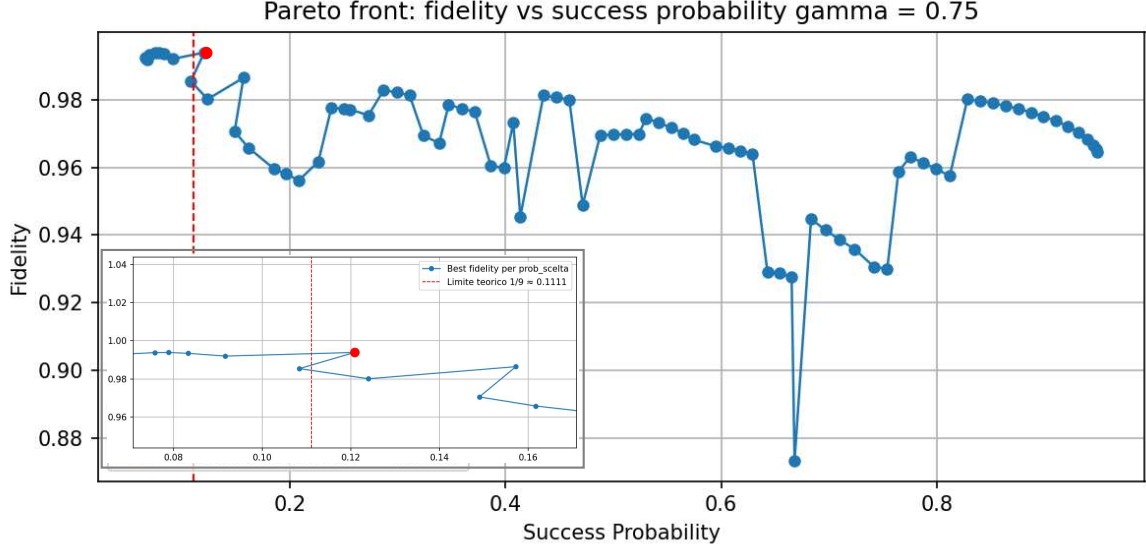


Figure B.16: Optimized Pareto front for $\Gamma = 0.75$ ($\Gamma/J_{\max} = 0.107$). The inset zooms in on the region around the highest-fidelity point, highlighted in red.

with a residual cost $C(\theta^*) = 0.0123$. The corresponding gate matrix, projected onto the logical subspace, is

$$G_{4 \times 4} = \begin{pmatrix} -0.096 - 0.364i & 0.008 + 0.006i & 0 & 0 \\ 0.008 + 0.006i & -0.122 - 0.284i & 0 - 0.001i & 0 \\ 0 & 0 - 0.001i & -0.010 + 0.032i & -0.113 - 0.331i \\ 0 & 0 & -0.113 - 0.331i & -0.023 + 0.020i \end{pmatrix}, \quad (\text{B.17})$$

and its normalized, phase-corrected version reads

$$\tilde{G} = \begin{pmatrix} 0.542 & 0.014 & 0.001 & 0 \\ 0.014 & 0.444 & 0.001 & 0.001 \\ 0.001 & 0.001 & 0.049 & 0.503 \\ 0 & 0.001 & 0.503 & 0.044 \end{pmatrix}, \quad (\text{B.18})$$

where entries below 10^{-3} have been set to zero. The block structure and the near-ideal diagonal values (0.503, 0.542, 0.444) are consistent with the high fidelity $\mathcal{F} = 0.9939$. Figures B.18 and B.19 show the time evolution of the photon density and its value at $T = 1$, for each computational basis input state. As in the $\Gamma = 0.5$ case, the target block no longer shows the clean two-site oscillation of the linear regime, with population spreading across a wider set of sites during the evolution before recombining close to the logical output at $T = 1$. Figure B.17 shows the corresponding normalized, dephased gate $\tilde{G}$.

Lowest-fidelity point One particular point along this front is worth examining in detail, not because it is the overall minimum, but because its underlying structure is qualitatively different from every other configuration discussed so far: $p_{\text{target}} = 0.6830$ ($p_{\text{succ}} = 0.6680$, $\mathcal{F} = 0.8734$). It is worth noting that, in addition to being the absolute

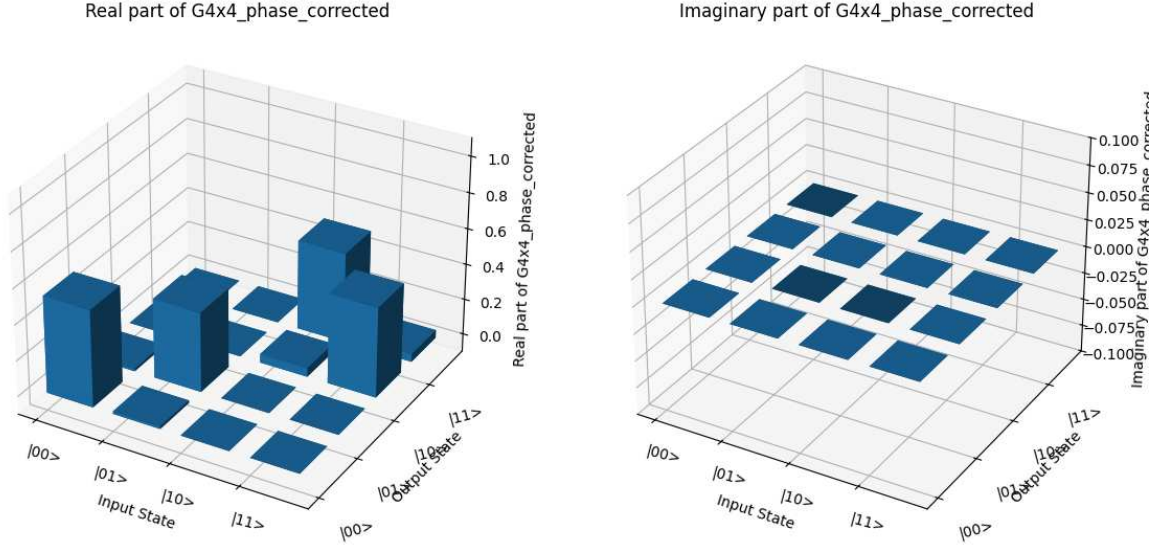


Figure B.17: Real and imaginary parts of the normalized, dephased logical gate $\tilde{G}$ at the highest-fidelity point for $\Gamma = 0.75$ ($p_{\text{succ}} = 0.1209$, $\mathcal{F} = 0.9939$).

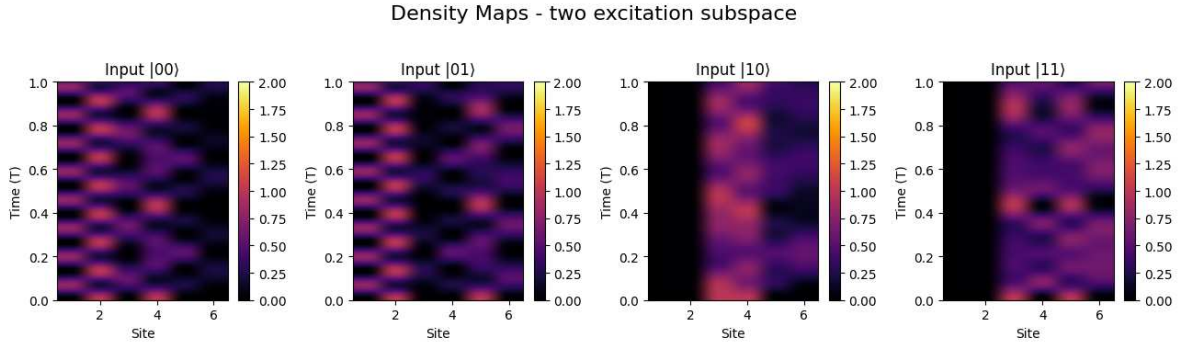


Figure B.18: Time evolution of the photon density across the six waveguides, from $T = 0$ to $T = 1$, for each computational basis input state, at the highest-fidelity point for $\Gamma = 0.75$ ($p_{\text{succ}} = 0.1209$, $\mathcal{F} = 0.9939$).

minimum of the Pareto front, it is also significantly lower than the surrounding points on either side of it ($\mathcal{F} \approx 0.93$). Normalizing the optimized parameters by $J_{\text{max}}^{\text{loc}} = |J_{45}^*| \approx 5.785 \pi^{-1}$ reveals a decoupling pattern different from every other configuration discussed so far: here it is $J_{12}/J_{\text{max}}^{\text{loc}} \approx 7 \cdot 10^{-9}$ and $J_{56}/J_{\text{max}}^{\text{loc}} \approx 7 \cdot 10^{-7}$ that vanish, rather than J_{23} . Physically, the two ancillary sites $|a_1\rangle$ and $|a_2\rangle$ become fully decoupled from the rest of the chain, isolating the four inner sites – the control and target qubits – as a self-contained four-site subsystem, connected internally by J_{23} , J_{34} , and J_{45} , all now non-negligible.

Figures B.21 and B.22 confirm this picture: the photon density is confined to sites 2–5 throughout the evolution, with no leakage onto the ancilla sites 1 and 6 at any time. However, unlike the tail of the linear case, this decoupling does not translate into a simple product structure. Since J_{23} , J_{34} , and J_{45} are all active simultaneously, population is free to hop across the entire four-site block, and the resulting dynamics visibly mixes the control and target degrees of freedom, rather than confining each to an isolated pair of sites. This is directly reflected in the normalized, phase-corrected

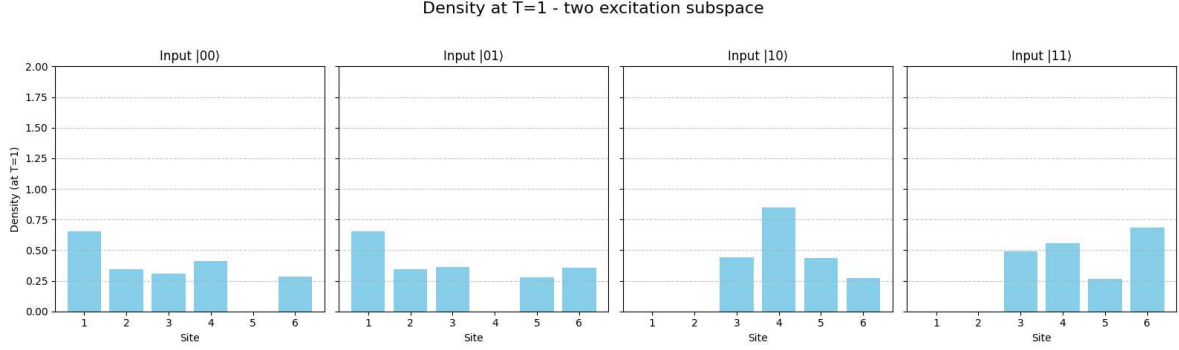


Figure B.19: Photon density at $T = 1$ across the six waveguides, for each computational basis input state, at the highest-fidelity point for $\Gamma = 0.75$ ($p_{\text{succ}} = 0.1209$, $\mathcal{F} = 0.9939$).

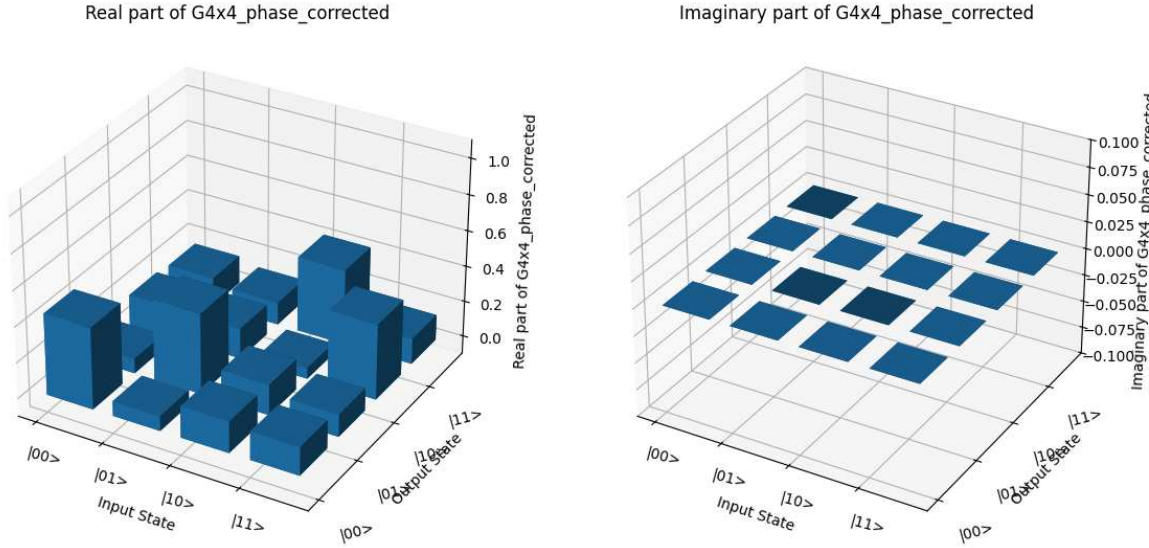


Figure B.20: Real and imaginary parts of the normalized, de-phased logical gate $\tilde{G}$ at the lowest-fidelity point of this Pareto front ($p_{\text{succ}} = 0.6680$, $\mathcal{F} = 0.8734$), showing the absence of the block-diagonal structure observed elsewhere in this chapter.

gate matrix,

$$\tilde{G} = \begin{pmatrix} 0.455 & 0.082 & 0.174 & 0.157 \\ 0.082 & 0.450 & 0.169 & 0.125 \\ 0.174 & 0.169 & 0.062 & 0.421 \\ 0.157 & 0.125 & 0.421 & 0.143 \end{pmatrix}, \quad (\text{B.19})$$

shown in Fig. B.20. In sharp contrast to every other configuration examined in this chapter, $\tilde{G}$ has no block-diagonal structure at all: the off-diagonal entries coupling the two qubit subspaces are comparable in magnitude to the diagonal blocks themselves, indicating that the gate substantially mix all four logical states into one another, rather than factorizing into two independent single-qubit operations as in the product-gate floor discussed in the linear case.

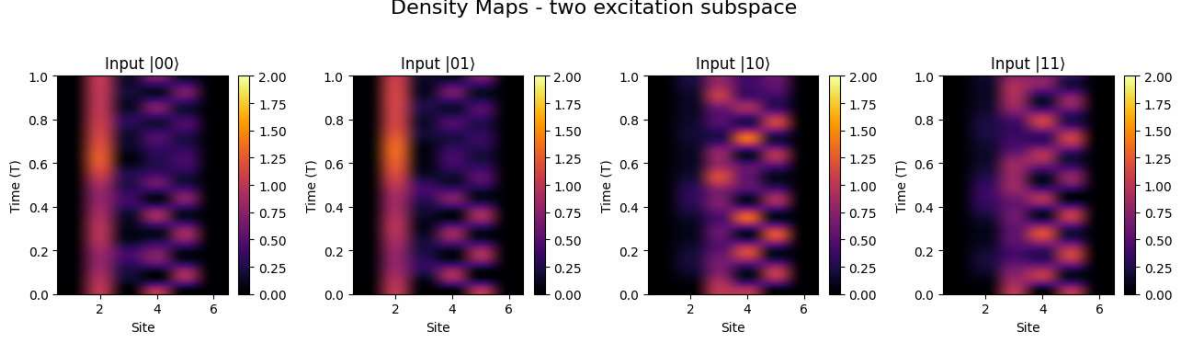


Figure B.21: Time evolution of the photon density across the six waveguides, from $T = 0$ to $T = 1$, for each computational basis input state, at the lowest-fidelity point of this Pareto front ($p_{\text{succ}} = 0.6680$, $\mathcal{F} = 0.8734$).

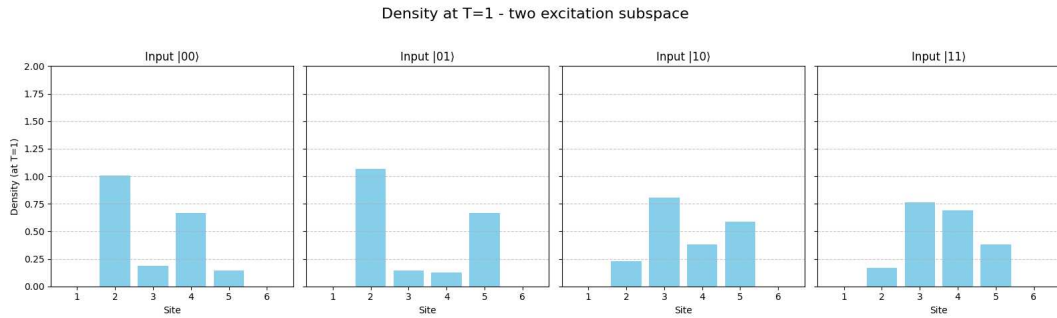


Figure B.22: Photon density at $T = 1$ across the six waveguides, for each computational basis input state, at the lowest-fidelity point of this Pareto front ($p_{\text{succ}} = 0.6680$, $\mathcal{F} = 0.8734$).

Bibliography

- [1] Michael A. Nielsen and Isaac L. Chuang. *Quantum Computation and Quantum Information: 10th Anniversary Edition*. Cambridge University Press, 2010.
- [2] David P. DiVincenzo. “The Physical Implementation of Quantum Computation”. In: *Fortschritte der Physik* 48.9-11 (2000), pp. 771–783. ISSN: 1521-3978. DOI: [10.1002/1521-3978\(200009\)48:9/11<771::aid-prop771>3.0.co;2-e](https://doi.org/10.1002/1521-3978(200009)48:9/11<771::aid-prop771>3.0.co;2-e).
- [3] Julio Gea-Banacloche. “Impossibility of large phase shifts via the giant Kerr effect with single-photon wave packets”. In: *Physical Review A* 81.4 (2010), p. 043823. DOI: [10.1103/PhysRevA.81.043823](https://doi.org/10.1103/PhysRevA.81.043823).
- [4] Jeffrey H. Shapiro. “Single-photon Kerr nonlinearities do not help quantum computation”. In: *Physical Review A* 73.6 (2006), p. 062305. DOI: [10.1103/PhysRevA.73.062305](https://doi.org/10.1103/PhysRevA.73.062305).
- [5] E. Knill, R. Laflamme, and G. J. Milburn. “A scheme for efficient quantum computation with linear optics”. In: *Nature* 409 (2001), pp. 46–52. DOI: [10.1038/35051009](https://doi.org/10.1038/35051009).
- [6] E. Knill. “Bounds on the Probability of Success of Postselected Nonlinear Sign Shifts Implemented with Linear Optics”. In: *Physical Review A* 68.6 (Dec. 2003), p. 064303. DOI: [10.1103/PhysRevA.68.064303](https://doi.org/10.1103/PhysRevA.68.064303). arXiv: [quant-ph/0307015](https://arxiv.org/abs/quant-ph/0307015).
- [7] Stefan Scheel and Norbert Lütkenhaus. “Upper bounds on success probabilities in linear optics”. In: *New Journal of Physics* 6 (2004), p. 51. DOI: [10.1088/1367-2630/6/1/051](https://doi.org/10.1088/1367-2630/6/1/051).
- [8] J. Eisert. “Optimizing Linear Optics Quantum Gates”. In: *Physical Review Letters* 95.4 (2005), p. 040502. DOI: [10.1103/PhysRevLett.95.040502](https://doi.org/10.1103/PhysRevLett.95.040502).
- [9] T. B. Pittman, B. C. Jacobs, and J. D. Franson. “Probabilistic quantum logic operations using polarizing beam splitters”. In: *Physical Review A* 64.6 (2001), p. 062311. DOI: [10.1103/PhysRevA.64.062311](https://doi.org/10.1103/PhysRevA.64.062311).
- [10] T. C. Ralph et al. “Linear optical controlled-NOT gate in the coincidence basis”. In: *Physical Review A* 65.6 (2002), p. 062324. DOI: [10.1103/PhysRevA.65.062324](https://doi.org/10.1103/PhysRevA.65.062324).
- [11] T. B. Pittman et al. “Experimental controlled-NOT logic gate for single photons in the coincidence basis”. In: *Physical Review A* 68.3 (2003), p. 032316. DOI: [10.1103/PhysRevA.68.032316](https://doi.org/10.1103/PhysRevA.68.032316).
- [12] Sara Gasparoni et al. “Realization of a Photonic Controlled-NOT Gate Sufficient for Quantum Computation”. In: *Physical Review Letters* 93.2 (2004), p. 020504. DOI: [10.1103/PhysRevLett.93.020504](https://doi.org/10.1103/PhysRevLett.93.020504).

- [13] J. L. O’Brien et al. “Demonstration of an all-optical quantum controlled-NOT gate”. In: *Nature* 426 (2003), pp. 264–267. DOI: [10.1038/nature02054](https://doi.org/10.1038/nature02054).
- [14] Alberto Politi et al. “Silica-on-Silicon Waveguide Quantum Circuits”. In: *Science* 320.5876 (2008), pp. 646–649. DOI: [10.1126/science.1155441](https://doi.org/10.1126/science.1155441).
- [15] Y. Lahini et al. “Quantum logic using correlated one-dimensional quantum walks”. In: *npj Quantum Information* 4 (2018), p. 2. DOI: [10.1038/s41534-017-0050-2](https://doi.org/10.1038/s41534-017-0050-2).
- [16] Robert J. Chapman et al. “Quantum logical controlled-NOT gate in a lithium niobate-on-insulator photonic quantum walk”. In: *Quantum Science and Technology* 9.1 (2024), p. 015016. DOI: [10.1088/2058-9565/ad0a48](https://doi.org/10.1088/2058-9565/ad0a48).
- [17] F. Scala, D. Nigro, and D. Gerace. “Deterministic entangling gates with nonlinear quantum photonic interferometers”. In: *Communications Physics* 7 (2024), p. 96. DOI: [10.1038/s42005-024-01610-z](https://doi.org/10.1038/s42005-024-01610-z).
- [18] Gregory R. Steinbrecher et al. “Quantum optical neural networks”. In: *npj Quantum Information* 5 (2019), p. 60. DOI: [10.1038/s41534-019-0174-7](https://doi.org/10.1038/s41534-019-0174-7).
- [19] Jacob Ewaniuk et al. “Imperfect Quantum Photonic Neural Networks”. In: *Advanced Quantum Technologies* 6.3 (2023), p. 2200125. DOI: [10.1002/qute.202200125](https://doi.org/10.1002/qute.202200125).
- [20] H. Edlbauer et al. “An 11-qubit atom processor in silicon”. In: *Nature* 648 (2025), pp. 569–575. DOI: [10.1038/s41586-025-09827-w](https://doi.org/10.1038/s41586-025-09827-w).
- [21] P. Steinacker et al. “Industry-compatible silicon spin-qubit unit cells exceeding 99% fidelity”. In: *Nature* 646 (2025), pp. 81–87. DOI: [10.1038/s41586-025-09531-9](https://doi.org/10.1038/s41586-025-09531-9).
- [22] Michael Foss-Feig et al. “Progress in Trapped-Ion Quantum Simulation”. In: *Annual Review of Condensed Matter Physics* 16 (2025), pp. 145–172. DOI: [10.1146/annurev-conmatphys-032822-045619](https://doi.org/10.1146/annurev-conmatphys-032822-045619).
- [23] M. P. Bland et al. “Millisecond lifetimes and coherence times in 2D transmon qubits”. In: *Nature* 647 (2025), pp. 343–348. DOI: [10.1038/s41586-025-09687-4](https://doi.org/10.1038/s41586-025-09687-4).
- [24] Nicolas Maring et al. “A versatile single-photon-based quantum computing platform”. In: *Nature Photonics* 18 (2024), pp. 603–609. DOI: [10.1038/s41566-024-01403-4](https://doi.org/10.1038/s41566-024-01403-4).
- [25] Pieter Kok et al. “Linear optical quantum computing with photonic qubits”. In: *Reviews of Modern Physics* 79.1 (2007), pp. 135–174. DOI: [10.1103/RevModPhys.79.135](https://doi.org/10.1103/RevModPhys.79.135).
- [26] F. Scala. “Leveraging Overparametrization: From Quantum Optimal Control to Quantum Machine Learning”. PhD thesis. Università degli Studi di Pavia, 2025.
- [27] Christophe Couteau. “Quantum computing using photons”. In: *The European Physical Journal A* 61 (2025), p. 75. DOI: [10.1140/epja/s10050-025-01517-5](https://doi.org/10.1140/epja/s10050-025-01517-5).
- [28] C. K. Hong, Z. Y. Ou, and L. Mandel. “Measurement of subpicosecond time intervals between two photons by interference”. In: *Phys. Rev. Lett.* 59 (18 Nov. 1987), pp. 2044–2046. DOI: [10.1103/PhysRevLett.59.2044](https://doi.org/10.1103/PhysRevLett.59.2044). URL: <https://link.aps.org/doi/10.1103/PhysRevLett.59.2044>.

- [29] Z.Y. Ou, C.K. Hong, and L. Mandel. “Relation between input and output states for a beam splitter”. In: *Optics Communications* 63.2 (1987), pp. 118–122. ISSN: 0030-4018. DOI: [https://doi.org/10.1016/0030-4018\(87\)90271-9](https://doi.org/10.1016/0030-4018(87)90271-9). URL: <https://www.sciencedirect.com/science/article/pii/0030401887902719>.
- [30] I. Aharonovich, K.B. Crozier, and D. Neshev. “Programmable integrated quantum photonics”. In: *Nature Photonics* (2026).
- [31] Peter P. Rohde, Timothy C. Ralph, and Michael A. Nielsen. “Optimal photons for quantum-information processing”. In: *Phys. Rev. A* 72 (5 Nov. 2005), p. 052332. DOI: [10.1103/PhysRevA.72.052332](https://doi.org/10.1103/PhysRevA.72.052332). URL: <https://link.aps.org/doi/10.1103/PhysRevA.72.052332>.
- [32] Niko Nikolay et al. “Very Large and Reversible Stark-Shift Tuning of Single Emitters in Layered Hexagonal Boron Nitride”. In: *Physical Review Applied* 11.4 (Apr. 2019). ISSN: 2331-7019. DOI: [10.1103/physrevapplied.11.041001](https://doi.org/10.1103/physrevapplied.11.041001). URL: <http://dx.doi.org/10.1103/PhysRevApplied.11.041001>.
- [33] David González-Andrade et al. “Ultra-broadband nanophotonic phase shifter based on subwavelength metamaterial waveguides”. In: *Photon. Res.* 8.3 (Mar. 2020), pp. 359–367. DOI: [10.1364/PRJ.373223](https://doi.org/10.1364/PRJ.373223). URL: <https://opg.optica.org/prj/abstract.cfm?URI=prj-8-3-359>.
- [34] Haidong Liang et al. “Tunable polarization entangled photon-pair source in rhombohedral boron nitride”. In: *Science Advances* 11.4 (2025), eadt3710. DOI: [10.1126/sciadv.adt3710](https://doi.org/10.1126/sciadv.adt3710). URL: <https://www.science.org/doi/abs/10.1126/sciadv.adt3710>.
- [35] Yaping Yang et al. “In situ manipulation of van der Waals heterostructures for twistronics”. In: *Science Advances* 6.49 (2020), eabd3655. DOI: [10.1126/sciadv.abd3655](https://doi.org/10.1126/sciadv.abd3655). URL: <https://www.science.org/doi/abs/10.1126/sciadv.abd3655>.
- [36] Manuel Erhard, Mario Krenn, and Anton Zeilinger. “Advances in high-dimensional quantum entanglement”. In: *Nature Reviews Physics* 2.7 (2020), pp. 365–381. DOI: [10.1038/s42254-020-0193-5](https://doi.org/10.1038/s42254-020-0193-5). URL: <https://doi.org/10.1038/s42254-020-0193-5>.
- [37] Chao He, Yijie Shen, and Andrew Forbes. “Towards higher-dimensional structured light”. In: *Light: Science & Applications* 11.1 (2022), p. 205. ISSN: 2047-7538. DOI: [10.1038/s41377-022-00897-3](https://doi.org/10.1038/s41377-022-00897-3). URL: <https://doi.org/10.1038/s41377-022-00897-3>.
- [38] Johannes Büttow et al. “Photonic integrated processor for structured light detection and distinction”. In: *Communications Physics* 6.1 (2023), p. 369. ISSN: 2399-3650. DOI: [10.1038/s42005-023-01489-2](https://doi.org/10.1038/s42005-023-01489-2). URL: <https://doi.org/10.1038/s42005-023-01489-2>.
- [39] Yun Zheng et al. “Multichip multidimensional quantum networks with entanglement retrievability”. In: *Science* 381.6654 (2023), pp. 221–226. DOI: [10.1126/science.adg9210](https://doi.org/10.1126/science.adg9210). URL: <https://www.science.org/doi/abs/10.1126/science.adg9210>.
- [40] Yaowen Hu et al. “On-chip electro-optic frequency shifters and beam splitters”. In: *Nature* 599.7886 (2021), pp. 587–593. DOI: [10.1038/s41586-021-03999-x](https://doi.org/10.1038/s41586-021-03999-x).

- [41] Chengli Wang et al. “Ultrabroadband thin-film lithium tantalate modulator for high-speed communications”. In: *Optica* 11.12 (Dec. 2024), pp. 1614–1620. DOI: [10.1364/OPTICA.537730](https://doi.org/10.1364/OPTICA.537730). URL: <https://opg.optica.org/optica/abstract.cfm?URI=optica-11-12-1614>.
- [42] Dong Kim et al. “Programmable photonic arrays based on microelectromechanical elements with femtowatt-level standby power consumption”. In: *Nature Photonics* 17 (Nov. 2023). DOI: [10.1038/s41566-023-01327-5](https://doi.org/10.1038/s41566-023-01327-5).
- [43] Rui Chen et al. “Opportunities and Challenges for Large-Scale Phase-Change Material Integrated Electro-Photonics”. In: *ACS Photonics* 9.10 (2022), pp. 3181–3195. DOI: [10.1021/acsp Photonics.2c00976](https://doi.org/10.1021/acsp Photonics.2c00976). URL: <https://doi.org/10.1021/acsp Photonics.2c00976>.
- [44] Shengping Liu et al. “Thermo-optic phase shifters based on silicon-on-insulator platform: state-of-the-art and a review”. In: *Frontiers of Optoelectronics* 15.1 (2022), p. 9. DOI: [10.1007/s12200-022-00012-9](https://doi.org/10.1007/s12200-022-00012-9).
- [45] Daryl Achilles et al. “Fiber-assisted detection with photon number resolution”. In: *Opt. Lett.* 28.23 (Dec. 2003), pp. 2387–2389. DOI: [10.1364/OL.28.002387](https://doi.org/10.1364/OL.28.002387). URL: <https://opg.optica.org/ol/abstract.cfm?URI=ol-28-23-2387>.
- [46] M. J. Fitch et al. “Photon-number resolution using time-multiplexed single-photon detectors”. In: *Phys. Rev. A* 68 (4 Oct. 2003), p. 043814. DOI: [10.1103/PhysRevA.68.043814](https://doi.org/10.1103/PhysRevA.68.043814). URL: <https://link.aps.org/doi/10.1103/PhysRevA.68.043814>.
- [47] Jacques Carolan et al. “Universal linear optics”. In: *Science* 349.6249 (2015), pp. 711–716. DOI: [10.1126/science.aab3642](https://doi.org/10.1126/science.aab3642).
- [48] Robert H. Hadfield. “Single-photon detectors for optical quantum information applications”. In: *Nature Photonics* 3.12 (2009), pp. 696–705. DOI: [10.1038/nphoton.2009.230](https://doi.org/10.1038/nphoton.2009.230).
- [49] François-Marie Le Régent. “Awesome Quantum Computing Experiments: Benchmarking Experimental Progress Towards Fault-Tolerant Quantum Computation”. In: (2025). arXiv: [2507.03678](https://arxiv.org/abs/2507.03678) [quant-ph]. URL: <https://arxiv.org/abs/2507.03678>.
- [50] Colin D. Bruzewicz et al. “Trapped-ion quantum computing: Progress and challenges”. In: *Applied Physics Reviews* 6.2 (May 2019). ISSN: 1931-9401. DOI: [10.1063/1.5088164](https://doi.org/10.1063/1.5088164). URL: <http://dx.doi.org/10.1063/1.5088164>.
- [51] Kenneth R. Brown et al. *Materials Challenges for Trapped-Ion Quantum Computers*. 2020. arXiv: [2009.00568](https://arxiv.org/abs/2009.00568) [quant-ph]. URL: <https://arxiv.org/abs/2009.00568>.
- [52] Karen Wintersperger et al. “Neutral atom quantum computing hardware: performance and end-user perspective”. In: *EPJ Quantum Technology* 10.1 (Aug. 2023). ISSN: 2196-0763. DOI: [10.1140/epjqt/s40507-023-00190-1](https://doi.org/10.1140/epjqt/s40507-023-00190-1). URL: <http://dx.doi.org/10.1140/epjqt/s40507-023-00190-1>.
- [53] Morten Kjaergaard et al. “Superconducting Qubits: Current State of Play”. In: *Annual Review of Condensed Matter Physics* 11.1 (Mar. 2020), pp. 369–395. ISSN: 1947-5462. DOI: [10.1146/annurev-conmatphys-031119-050605](https://doi.org/10.1146/annurev-conmatphys-031119-050605). URL: <http://dx.doi.org/10.1146/annurev-conmatphys-031119-050605>.

- [54] Sébastien Pezzagna and J. Meijer. “Quantum computer based on color centers in diamond”. In: *Applied Physics Reviews* 8 (Mar. 2021), p. 011308. DOI: [10.1063/5.0007444](https://doi.org/10.1063/5.0007444).
- [55] Arne Laucht et al. “Roadmap on quantum nanotechnologies”. In: *Nanotechnology* 32.16 (Feb. 2021), p. 162003. ISSN: 1361-6528. DOI: [10.1088/1361-6528/abb333](https://doi.org/10.1088/1361-6528/abb333). URL: <http://dx.doi.org/10.1088/1361-6528/abb333>.
- [56] François-Marie Le Régent. *Awesome Quantum Computing Experiments: Benchmarking Experimental Progress Towards Fault-Tolerant Quantum Computation*. 2025. DOI: [10.48550/arXiv.2507.03678](https://doi.org/10.48550/arXiv.2507.03678). arXiv: [2507.03678](https://arxiv.org/abs/2507.03678) [quant-ph].
- [57] Robert Raussendorf and Hans J. Briegel. “A One-Way Quantum Computer”. In: *Phys. Rev. Lett.* 86 (22 May 2001), pp. 5188–5191. DOI: [10.1103/PhysRevLett.86.5188](https://doi.org/10.1103/PhysRevLett.86.5188). URL: <https://link.aps.org/doi/10.1103/PhysRevLett.86.5188>.
- [58] Robert Raussendorf, Daniel E. Browne, and Hans J. Briegel. “Measurement-based quantum computation on cluster states”. In: *Physical Review A* 68.2 (Aug. 2003). ISSN: 1094-1622. DOI: [10.1103/physreva.68.022312](https://doi.org/10.1103/physreva.68.022312). URL: <http://dx.doi.org/10.1103/PhysRevA.68.022312>.
- [59] Michael A. Nielsen. “Optical Quantum Computation Using Cluster States”. In: *Physical Review Letters* 93.4 (July 2004). ISSN: 1079-7114. DOI: [10.1103/physrevlett.93.040503](https://doi.org/10.1103/physrevlett.93.040503). URL: <http://dx.doi.org/10.1103/PhysRevLett.93.040503>.
- [60] E. Knill. “Quantum Gates Using Linear Optics and Postselection”. In: *Physical Review A* 66.5 (Nov. 2002), p. 052306. DOI: [10.1103/PhysRevA.66.052306](https://doi.org/10.1103/PhysRevA.66.052306). arXiv: [quant-ph/0110144](https://arxiv.org/abs/quant-ph/0110144).
- [61] S. Scheel et al. “Feed-forward and its role in conditional linear optical quantum dynamics”. In: *Phys. Rev. A* 73 (3 Mar. 2006), p. 034301. DOI: [10.1103/PhysRevA.73.034301](https://doi.org/10.1103/PhysRevA.73.034301). URL: <https://link.aps.org/doi/10.1103/PhysRevA.73.034301>.
- [62] Daniel Gottesman and Isaac L. Chuang. “Demonstrating the viability of universal quantum computation using teleportation and single-qubit operations”. In: *Nature* 402.6760 (Nov. 1999), pp. 390–393. DOI: [10.1038/46503](https://doi.org/10.1038/46503). arXiv: [quant-ph/9908010](https://arxiv.org/abs/quant-ph/9908010) [quant-ph].
- [63] Charles H. Bennett et al. “Teleporting an unknown quantum state via dual classical and Einstein-Podolsky-Rosen channels”. In: *Phys. Rev. Lett.* 70 (13 Mar. 1993), pp. 1895–1899. DOI: [10.1103/PhysRevLett.70.1895](https://doi.org/10.1103/PhysRevLett.70.1895). URL: <https://link.aps.org/doi/10.1103/PhysRevLett.70.1895>.
- [64] Y. Aharonov, L. Davidovich, and N. Zagury. “Quantum random walks”. In: *Phys. Rev. A* 48 (2 Aug. 1993), pp. 1687–1690. DOI: [10.1103/PhysRevA.48.1687](https://doi.org/10.1103/PhysRevA.48.1687). URL: <https://link.aps.org/doi/10.1103/PhysRevA.48.1687>.
- [65] J Kempe. “Quantum random walks: An introductory overview”. In: *Contemporary Physics* 44.4 (July 2003), pp. 307–327. ISSN: 1366-5812. DOI: [10.1080/00107151000110776](https://doi.org/10.1080/00107151000110776). URL: <http://dx.doi.org/10.1080/00107151000110776>.
- [66] Hagai B. Perets et al. “Realization of Quantum Walks with Negligible Decoherence in Waveguide Lattices”. In: *Physical Review Letters* 100.17 (2008), p. 170506. DOI: [10.1103/PhysRevLett.100.170506](https://doi.org/10.1103/PhysRevLett.100.170506).

- [67] Salvador Elías Venegas-Andraca. “Quantum walks: a comprehensive review”. In: *Quantum Information Processing* 11.5 (July 2012), pp. 1015–1106. ISSN: 1573-1332. DOI: [10.1007/s11128-012-0432-5](https://doi.org/10.1007/s11128-012-0432-5). URL: <http://dx.doi.org/10.1007/s11128-012-0432-5>.
- [68] C. M. Chandrashekar, Sandeep K. Goyal, and Subhashish Banerjee. “Entanglement Generation in Spatially Separated Systems Using Quantum Walk”. In: *Journal of Quantum Information Science* 2.2 (2012), pp. 15–22. DOI: [10.4236/jqis.2012.22004](https://doi.org/10.4236/jqis.2012.22004).
- [69] D. Nigro et al. “Integrated quantum polariton interferometry”. In: *Communications Physics* 5 (2022), p. 34. DOI: [10.1038/s42005-022-00810-9](https://doi.org/10.1038/s42005-022-00810-9).
- [70] D.F. Walls and G.J. Milburn. *Quantum Optics*. Springer Study Edition. Springer Berlin Heidelberg, 1995. ISBN: 9783540588313. URL: <https://books.google.it/books?id=Vx5RAAAAMAAJ>.
- [71] Anton Frisk Kockum et al. “Deterministic quantum nonlinear optics with single atoms and virtual photons”. In: *Phys. Rev. A* 95 (6 June 2017), p. 063849. DOI: [10.1103/PhysRevA.95.063849](https://doi.org/10.1103/PhysRevA.95.063849). URL: <https://link.aps.org/doi/10.1103/PhysRevA.95.063849>.
- [72] Xiaoguang Wang, Chang-Shui Yu, and X. X. Yi. “An alternative quantum fidelity for mixed states of qudits”. In: *arXiv preprint* (2008). arXiv: [0807.1781](https://arxiv.org/abs/0807.1781) [quant-ph].
- [73] A. Matthew Smith et al. “Imperfect linear optical photonic gates with number-resolving photodetection”. In: *Physical Review A* 84 (2011), p. 032341. DOI: [10.1103/PhysRevA.84.032341](https://doi.org/10.1103/PhysRevA.84.032341).
- [74] Abhinav Kandala et al. “Error mitigation extends the computational reach of a noisy quantum processor”. In: *Nature* 567.7749 (2019), pp. 491–495. DOI: [10.1038/s41586-019-1040-7](https://doi.org/10.1038/s41586-019-1040-7).
- [75] I. L. Chuang and Y. Yamamoto. “Simple quantum computer”. In: *Physical Review A* 52 (1995), pp. 3489–3496. DOI: [10.1103/PhysRevA.52.3489](https://doi.org/10.1103/PhysRevA.52.3489).
- [76] Atharv Joshi, Kyungjoo Noh, and Yvonne Gao. “Quantum information processing with bosonic qubits in circuit QED”. In: *Quantum Science and Technology* 6 (Apr. 2021). DOI: [10.1088/2058-9565/abe989](https://doi.org/10.1088/2058-9565/abe989).
- [77] Michael A. Nielsen. “Cluster-state quantum computation”. In: *Reports on Mathematical Physics* 57.1 (Feb. 2006), pp. 147–161. ISSN: 0034-4877. DOI: [10.1016/S0034-4877\(06\)80014-5](https://doi.org/10.1016/S0034-4877(06)80014-5). URL: [http://dx.doi.org/10.1016/S0034-4877\(06\)80014-5](http://dx.doi.org/10.1016/S0034-4877(06)80014-5).
- [78] Masato Koashi, Takashi Yamamoto, and Nobuyuki Imoto. “Probabilistic manipulation of entangled photons”. In: *Physical Review A* 63.3 (2001), p. 030301. DOI: [10.1103/PhysRevA.63.030301](https://doi.org/10.1103/PhysRevA.63.030301).
- [79] Edward Farhi and Sam Gutmann. “Quantum computation and decision trees”. In: *Phys. Rev. A* 58 (2 Aug. 1998), pp. 915–928. DOI: [10.1103/PhysRevA.58.915](https://doi.org/10.1103/PhysRevA.58.915). URL: <https://link.aps.org/doi/10.1103/PhysRevA.58.915>.

- [80] W. Dür et al. “Quantum walks in optical lattices”. In: *Phys. Rev. A* 66 (5 Nov. 2002), p. 052319. DOI: [10.1103/PhysRevA.66.052319](https://doi.org/10.1103/PhysRevA.66.052319). URL: <https://link.aps.org/doi/10.1103/PhysRevA.66.052319>.
- [81] Alberto Peruzzo et al. “Quantum Walks of Correlated Photons”. In: *Science* 329.5998 (Sept. 2010), pp. 1500–1503. ISSN: 1095-9203. DOI: [10.1126/science.1193515](https://doi.org/10.1126/science.1193515). URL: <http://dx.doi.org/10.1126/science.1193515>.
- [82] J. J. Hopfield. “Neural networks and physical systems with emergent collective computational abilities”. In: *Proceedings of the National Academy of Sciences* 79 (1982), pp. 2554–2558. DOI: [10.1073/pnas.79.8.2554](https://doi.org/10.1073/pnas.79.8.2554).
- [83] Rahul Nandkishore and David A. Huse. “Many-body localization and thermalization in quantum statistical mechanics”. In: *Annual Review of Condensed Matter Physics* 6 (2015), pp. 15–38. DOI: [10.1146/annurev-conmatphys-031214-014726](https://doi.org/10.1146/annurev-conmatphys-031214-014726).
- [84] Andrea Crespi et al. “Anderson localization of entangled photons in an integrated quantum walk”. In: *Nature Photonics* 7 (2013), pp. 322–328. DOI: [10.1038/nphoton.2013.26](https://doi.org/10.1038/nphoton.2013.26).