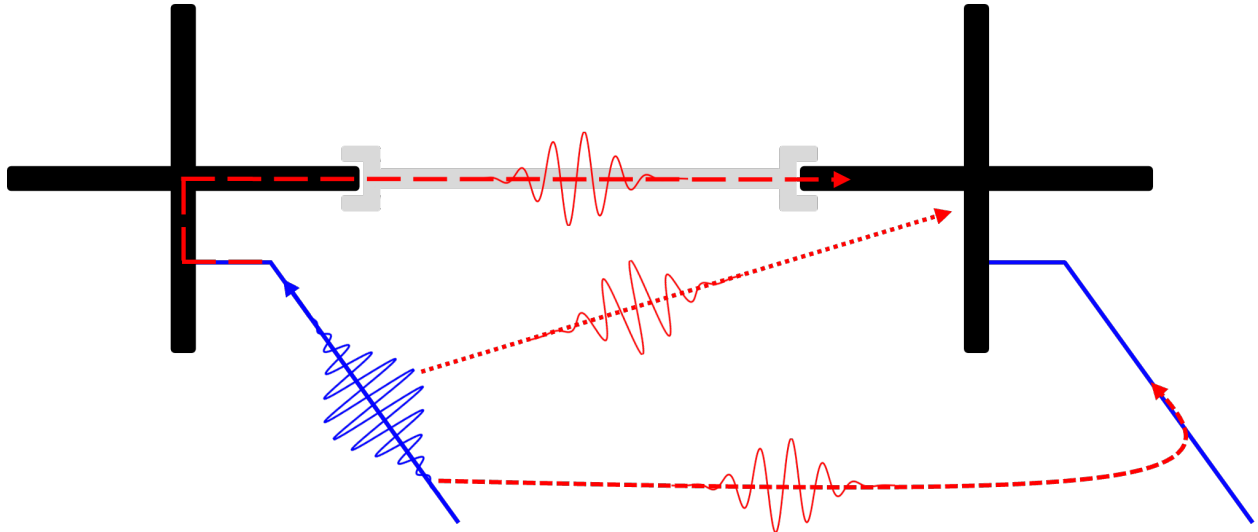




XY Crosstalk in a 25-Qubit Flip-Chip Superconducting Quantum Processor

Thesis for Erasmus Mundus Master of Science in Nanoscience and Nanotechnology

Specialisation: Quantum Computing

Promoter: Dr. Michele Faucci Giannelli, Chalmers University of Technology

Co-promoter: Prof. Bart Sorée, KU Leuven

Examiner: Dr. Axel Martin Eriksson, Chalmers University of Technology

HONG QUAN TRINH

DEPARTMENT OF MICROT TECHNOLOGY AND NANOSCIENCE

CHALMERS UNIVERSITY OF TECHNOLOGY & KU LEUVEN

Gothenburg, Sweden 2026

www.chalmers.se • www.kuleuven.be

MASTER'S THESIS 2026

XY Crosstalk in a 25-Qubit Flip-Chip Superconducting Quantum Processor

HONG QUAN TRINH



Department of Microtechnology and Nanoscience
Division of Quantum Technology
Quantum Computing (QC2) Group
CHALMERS UNIVERSITY OF TECHNOLOGY
Gothenburg, Sweden 2026

XY Crosstalk in a 25-Qubit Flip-Chip Superconducting Quantum Processor
HONG QUAN TRINH

© HONG QUAN TRINH, 2026.

Promoter: Dr. Michele Faucci Giannelli, Chalmers University of Technology
Co-promoter: Prof. Bart Sorée, KU Leuven
Examiner/External referee: Dr. Axel Martin Eriksson,
Chalmers University of Technology

Master's Thesis 2026
Department of Microtechnology and Nanoscience
Division of Quantum Technology
Quantum Computing (QC2) Group
Chalmers University of Technology
SE-412 96 Gothenburg
Telephone +46 31 772 1000

Cover: XY crosstalk between two superconducting qubits. Typeset in L^AT_EX

Printed by Chalmers Reproservice
Gothenburg, Sweden 2026

Abstract

Superconducting circuits are among the most promising physical platforms for building a quantum computer. Scaling superconducting quantum processors to thousands and eventually millions of qubits requires not only high-fidelity quantum gates but also the ability to calibrate and execute them simultaneously. XY crosstalk, resulting from electromagnetic coupling between neighbouring qubits and control lines, is one of the main obstacles to this goal, as it introduces coherent errors and leakage that degrade gate performance during parallel operations. While several crosstalk characterisation and mitigation methods have been demonstrated in superconducting quantum processors, a systematic and scalable approach for integrating them into an automated calibration workflow remains a practical challenge, especially for fixed-frequency architectures, where qubit frequencies cannot be tuned away from one another. In this thesis, an experimental protocol has been developed to measure and compensate XY crosstalk at the signal level in Chalmers' 25-qubit flip-chip fixed-frequency transmon quantum processor. It is shown that the amplitude and phase of the complex XY crosstalk coefficient can be extracted separately for both $f_{01} - f_{01}$ and $f_{01} - f_{12}$ crosstalk, and that these coefficients can be used for effective active compensation to improve simultaneous gate performance. Specifically, randomised benchmarking shows that by applying the crosstalk compensation scheme, the average single-qubit gate error during simultaneous measurements can be reduced to the same level as the value obtained from isolated operations. The same improvement in leakage rate is also observed for qubit pairs affected by $f_{01} - f_{12}$ crosstalk. These results prove that crosstalk-induced coherent errors and leakage errors can be completely removed using calibration and control strategies and represent a step towards reliable parallel operations in large-scale quantum processors.

Keywords: quantum computing, superconducting quantum processor, transmon qubits, flip-chip, XY crosstalk, electromagnetic coupling, microwave, coherent errors, leakage errors, scaling, parallel operations.

Acknowledgements

First and foremost, I would like to thank my supervisor, Dr. Michele Faucci Giannelli, and the PI of the QC2 group, Dr. Giovanna Tancredi, for giving me the opportunity to do my thesis and learn so much about quantum computing in the group. Michele, thank you for your guidance, patience, ideas, and all the long discussions we had over the past months. Even though you were always busy, you somehow managed to find time for my questions, help me whenever I got stuck, and point me in the right direction, while still giving me the freedom I need to explore on my own and grow as an independent researcher. I deeply appreciate the chance to have worked closely with someone as knowledgeable and experienced as you.

I would also like to thank the amazing people in the QC2 group - Abdullah, Anuj, Axel, Emil, Ingrid, Job, Joel, Junior, Matteo, and everyone else - for all your help with the software and measurements, for the many discussions and ideas, and for making the lab such an enjoyable place to work. I learned a lot from all of you, and I am grateful to have been part of the group.

To my family and friends, thank you for always believing in me and being there for me throughout this thesis, in the past two years of my Master's degree, and long before that. Whether through encouragement, advice, or simply sending me funny Instagram reels from time to time, your support has meant more than you probably realise. I could not have made it this far without you.

Finally, I would like to extend my gratitude to the EMM-Nano programme consortium for awarding me the Erasmus Mundus Joint Master full scholarship, which has made the past two years possible. Studying in Europe has not always been so easy. It was a crazy ride full of chaos, struggles, stressful moments, and uncertainty; but it also brought me endless joy, unforgettable memories, incredible people, and valuable opportunities I could have never imagined. Looking back, I realise this journey has taught me far more than nanotechnology and quantum computing alone. It has taught me strength, maturity, independence, and resilience. For all of that, I am forever grateful.

Hong Quan Trinh, Göteborg, August 2026

List of Acronyms

- C-chip** control chip. 13, 14, 31, 32
cQED circuit quantum electrodynamics. 5, 11
CZ controlled-phase. 10, 38
- DRAG** derivative removal by adiabatic gate. 24, 66, 70, 76
- FTQC** fault-tolerant quantum computing. 2, 5, 13
- HEMT** high electron mobility transistor. 35, 37
- IF** intermediate frequency. 37
- JJ** Josephson junction. 6–9, 34
- LO** local oscillator. 36, 46
- NISQ** noisy intermediate-scale quantum. 2
- PCB** printed circuit board. 14, 16, 32
- Q-chip** quantum chip. 13, 14, 31, 32
QEC quantum error correction. 2, 3, 13, 15
QND quantum non-demolition. 5, 11
QOI quantity of interest. 38
- RB** randomised benchmarking. 11, 16, 18, 19, 38, 42, 46–51, 53, 55, 73, 76, 77, 80–84
- SPAM** state preparation and measurement. 73
SQUID superconducting quantum interference device. 6–9, 32
SSRO single-shot readout. 11, 46, 53, 55, 64, 76
- UBM** under-bump metallisation. 14

Parameters

$ \psi\rangle$	Qubit state
E_J	Josephson energy
E_C	Charging energy
$\hat{H}$	Hamiltonian
I_c	Critical current of the JJ
I_c^{SQUID}	Critical current of the SQUID
Φ_0	Flux quantum
$E_{01(12)}$	First (Second) transition energy of the qubit
$f_{01(12)}$	First (Second) transition frequency of the qubit
$\omega_{01(12)}$	First (Second) transition angular frequency of the qubit
α	Qubit anharmonicity
Δ	Frequency detuning
ω_c	Resonator frequency
$\hat{a}$	Annihilation operator
$\hat{a}^\dagger$	Creation operator
$\hat{\sigma}_+$	Raising operator
$\hat{\sigma}_-$	Lowering operator
$\hat{\sigma}_x, \hat{\sigma}_y, \hat{\sigma}_z$	Pauli X, Y, Z operators
g	Qubit-resonator coupling strength
χ	Dispersive shift
T_1	Relaxation time of the qubit
T_2	Total dephasing time of the qubit
T_2^*	Pure dephasing time of the qubit
T_2^{Echo}	Spin-echo dephasing time of the qubit
β	Bare crosstalk coefficient
$ \beta $	Crosstalk amplitude
φ	Crosstalk phase
φ^{canc}	Cancellation phase
Λ	Crosstalk level (in dB)
β^{eff}	Effective crosstalk coefficient
$ \tilde{\beta} $	Amplitude scaling factor for crosstalk phase measurement
C	Crosstalk matrix
Ω	Drive strength
A	Pulse amplitude
A^π	π pulse amplitude
τ	Pulse duration

θ	Rotation angle
$S(\Delta)$	Spectral overlap
β_{D}	DRAG parameter
γ	Unitless DRAG parameter
L	Leakage rate
S	Seepage rate
F_{avg}	Average Clifford gate fidelity
ε_{avg}	Average Clifford gate error
ε_{iso}	Average Clifford gate error obtained from isolated RB
ε_{sim}	Average Clifford gate error obtained from simultaneous RB

Contents

List of Acronyms	v
List of Parameters	vii
List of Figures	xi
1 Introduction	1
1.1 Quantum vs. Classical	2
1.2 Physical realisation of a quantum computer	4
1.3 Superconducting quantum computing	5
1.3.1 The JJ	6
1.3.2 From Cooper pair box to transmon qubit	7
1.3.3 Gate operations in superconducting circuits	10
1.3.4 Dispersive readout of superconducting qubits	11
1.3.5 Qubit decoherence	11
1.4 Scaling superconducting quantum processors	13
1.4.1 Motivation for scaling and 3D integration	13
1.4.2 Flip-chip integration	13
1.4.3 Frequency-multiplexed qubit control and readout	14
1.5 Crosstalk in a superconducting quantum processor	15
1.5.1 Classical crosstalk	15
1.5.1.1 XY crosstalk	16
1.5.2 Quantum crosstalk	17
1.6 In this Master's thesis	18
1.6.1 Scope & objectives	18
1.6.2 Thesis outline	18
2 XY Crosstalk: Theory & Literature Review	21
2.1 Mathematical formulation of XY crosstalk	21
2.1.1 Crosstalk coefficient & crosstalk matrix	21
2.1.2 The effect of frequency detuning & effective crosstalk coefficient	23
2.2 XY crosstalk measurement	24
2.2.1 Rabi-like measurement family	25
2.2.2 Ramsey-like & spin-echo measurement families	25
2.3 XY crosstalk compensation	28
3 Experimental Setup	31

3.1	The 25-qubit flip-chip quantum processor	31
3.2	Qubit frequency allocation scheme	31
3.3	Cryogenic measurement setup	34
3.4	The software stack	38
4	XY Crosstalk: Methods, Results & Discussion	39
4.1	XY crosstalk amplitude measurement	39
4.2	XY crosstalk phase measurement	42
4.3	XY crosstalk compensation	46
5	Conclusion & Future Outlook	53
5.1	Conclusion	53
5.2	Future outlook	53
A	Single-Qubit Gate Calibration & Characterisation	55
A.1	Readout calibration nodes	57
A.1.1	Resonator spectroscopy	57
A.1.2	Punch-out	57
A.1.3	State-discrimination readout frequency optimisation	58
A.1.4	State-discrimination readout amplitude optimisation	58
A.2	Qubit calibration nodes	65
A.2.1	Qubit spectroscopy	65
A.2.2	Rabi oscillations	65
A.2.3	Ramsey correction	65
A.2.4	DRAG parameter calibration	66
A.2.5	n-Rabi oscillations	66
A.3	Qubit characterisation nodes	72
A.3.1	T_1 measurement	72
A.3.2	T_2 measurement	72
A.3.3	T_2 Echo measurement	72
A.3.4	RB	73
A.4	Crosstalk measurement nodes	76
A.5	Example of a single-qubit gate calibration chain	76
B	Supplementary Materials	77
	Bibliography	85
	Code of Conduct & Transparency Statement on the Use of GenAI	89

List of Figures

1.1	Bloch sphere representing the state of a single qubit.	3
1.2	Illustration of a JJ.	7
1.3	The anharmonic energy potential of the transmon qubit.	8
1.4	Circuit diagrams of superconducting qubits.	9
1.5	The dependence of energy levels on the offset charge.	10
1.6	Dispersive readout of superconducting qubits.	12
1.7	Illustration of a two-qubit flip-chip superconducting quantum processor.	14
1.8	Possible XY crosstalk channels between two superconducting qubits.	17
2.1	XY crosstalk between two qubits.	22
2.2	A Rabi-like XY crosstalk measurement scheme.	26
2.3	A spin-echo XY crosstalk measurement scheme.	27
3.1	The 25-qubit flip-chip quantum processor.	32
3.2	Qubit frequency allocation scheme in the 25-qubit device.	33
3.3	Wiring diagram for the measurement setup at room and cryogenic temperatures.	35
3.4	Qblox clusters.	36
3.5	The experimental setup.	37
3.6	The software framework.	38
4.1	Pulse sequences for XY crosstalk measurement	40
4.2	Representative $f_{01} - f_{12}$ crosstalk amplitude results for qubits q07 and q09.	43
4.3	Representative $f_{01} - f_{12}$ crosstalk phase results for qubits q07 and q09.	45
4.4	RB results for qubit q09 without crosstalk compensation.	47
4.5	RB results for qubit q09 with crosstalk compensation.	48
4.6	Average single-qubit gate error matrix for qubits q07 and q09.	49
4.7	Leakage rate matrix for qubits q07 and q09.	50
4.8	Comparison of gate-error ratios with and without XY crosstalk compensation.	51
A.1	Single-qubit gate calibration node dependency tree.	56
A.2	Resonator spectroscopy for state $ 0\rangle$	60
A.3	Punch-out for state $ 0\rangle$	61
A.4	Readout spectroscopy for state $ 0\rangle$ after punch-out.	62
A.5	Three-state readout optimisation.	63

A.6	Three-state SSRO with optimised readout frequency and amplitude. .	64
A.7	Qubit spectroscopy for the $ 0\rangle - 1\rangle$ transition.	67
A.8	Amplitude-swept Rabi oscillations for the $ 0\rangle - 1\rangle$ transition.	68
A.9	Ramsey correction for the $ 0\rangle - 1\rangle$ transition.	69
A.10	DRAG parameter optimisation for the $ 0\rangle - 1\rangle$ transition.	70
A.11	n-Rabi oscillations for the $ 0\rangle - 1\rangle$ transition.	71
A.12	T_1 measurement.	74
A.13	T_2 Echo measurement.	75
B.1	Representative $f_{01} - f_{01}$ crosstalk amplitude results for qubits q11 and q15.	78
B.2	Representative $f_{01} - f_{01}$ crosstalk phase results for qubits q11 and q15.	79
B.3	RB results for qubit q07 without crosstalk compensation.	80
B.4	RB results for qubit q07 with crosstalk compensation.	81
B.5	Average single-qubit gate error matrix for qubits q11 and q13.	82
B.6	Average single-qubit gate error matrix for qubits q13 and q15.	83
B.7	Average single-qubit gate error matrix for qubits q11 and q15.	84

1

Introduction

Nature isn't classical, dammit, and if you want to make a simulation of Nature, you'd better make it quantum mechanical, and by golly it's a wonderful problem because it doesn't look so easy.

Richard Feynman, 1981

Forty-five years ago, in a famous speech titled “Simulating Physics with Computers”, Richard Feynman proposed a novel and more powerful type of computer - something fundamentally different from a Turing machine, that exploits the power of quantum mechanics to simulate complex systems in Nature [1]. Around the same period, Paul Benioff developed a quantum mechanical model of Turing machines and established the theoretical possibility of a reversible quantum computer, i.e., one that, in principle, does not dissipate any energy [2, 3].

The early work of Feynman, Benioff, and others have launched the quest for a new paradigm of computation over the following decades. Throughout the 1980s, foundational ideas began to take shape, most notably through the work of David Deutsch, who in 1985 formalised the idea of a universal quantum computer and demonstrated that quantum devices could perform computations beyond the reach of classical ones [4]. In Deutsch's toy model, a quantum computer can be viewed as a giant quantum interferometer, where probability amplitudes interfere with one another, allowing normally inaccessible internal information to manifest in measurable outcomes - a phenomenon often referred to as *quantum parallelism*, or *quantum interference*.

The field really gained significant momentum in the 1990s with the development of quantum algorithms that demonstrated potential computational advantages. In 1997, Daniel Simon showed an algorithm that exploits quantum interference to provide an exponential speedup over all known classical algorithms for a period-finding problem [5]. Building on Simon's idea, Peter Shor then proposed an efficient method for factorising large integers based on quantum Fourier transformation [6]. Another example is Grover's algorithm, which also uses quantum parallelism to achieve a quadratic speedup in searching an unsorted database [7]. These breakthroughs offered compelling evidence that quantum computers can outperform their classical counterparts in certain problems of practical importance, implying their applications in various fields such as computational chemistry and cryptography.

However, translating these theoretical advantages into practical devices has proven extraordinarily challenging. Quantum systems are inherently fragile: interactions with their surrounding environment result in the loss of quantum information, or *decoherence*, and hence operational errors that rapidly degrade computational accuracy. To address this, the theory of quantum error correction (QEC) has been established and is now one of the most active research directions in quantum computing. QEC enables the detection and correction of errors by encoding logical qubits into many physical qubits without directly measuring the encoded information. However, it comes at the cost of substantial overhead in the number of physical qubits required, which is in the order of hundreds of thousands to millions of qubits for practical applications.

In the absence of full error correction, current devices operate in the so-called noisy intermediate-scale quantum (NISQ) era, with only tens to hundreds of well-controlled qubits. The future goal is, therefore, to achieve fault-tolerant quantum computing (FTQC), where errors are actively detected and corrected throughout the computation, enabling long and reliable quantum circuits. With rapid progress in recent years, both in academia and industry, one can optimistically hope that that future with FTQC is not so far ahead at all. IBM, a leading player in the superconducting qubit platform, has outlined a roadmap towards large-scale FTQC systems, with a target to execute 100 million quantum gates on 200 qubits by 2029 [8]. Similar efforts by Google and others continue to push the boundaries of coherence, scalability, and error correction [9].

Forty-five years after Feynman’s great vision, though many remarkable advances have been made, the problem still, in his own words, *doesn’t look so easy* at all.

1.1 Quantum vs. Classical

Classically, information is stored and processed in bits, each of which can only take a definite and distinguishable value of either 0 or 1. Abstractly, this can be thought of as a light switch being either on or off, a vector pointing either up or down, or any physical system restricted to two possible discrete states.

In a quantum computer, information is instead encoded in *quantum bits* or *qubits*. The crucial difference is that a qubit can exist in a *superposition* of its basis states $|0\rangle$ and $|1\rangle$. One possible mathematical representation of a qubit state $|\psi\rangle$ is a vector in a complex Hilbert space spanned by its orthogonal basis states $|0\rangle$ and $|1\rangle$:

$$|\psi\rangle = \alpha |0\rangle + \beta |1\rangle, \quad |\alpha|^2 + |\beta|^2 = 1, \quad (1.1)$$

where α and β are complex probability amplitudes. Upon a projective measurement, however, the qubit state will probabilistically collapse into either $|0\rangle$ or $|1\rangle$ with probabilities $|\alpha|^2$ and $|\beta|^2$, respectively. This departure from classical determinism is one of the key features of quantum computation.

A useful way to visualise a single qubit state is via the Bloch sphere (Fig. 1.1), where any pure state $|\psi\rangle$ corresponds to an arbitrary unit vector on the surface of the sphere. In this picture, $|\psi\rangle$ can thus also be written in terms of two real angles θ and ϑ as

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

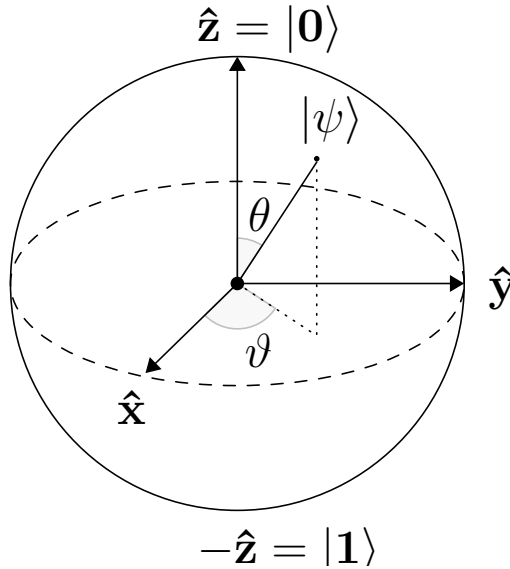


Figure 1.1: Bloch sphere representing the state of a single qubit.

with $0 \leq \theta \leq \pi$ and $0 \leq \vartheta < 2\pi$. From Eq. 1.2, it is clear that a single qubit, despite its complex form in Eq. 1.1, is fully described by only two real parameters¹. A quantum gate is a unitary operation that rotates the state vector by a specific angle about a specific axis on the Bloch sphere. It is thus used to manipulate the state of a qubit. It is also interesting to note that, within this representation, a classical bit can be viewed as just a special case of a qubit where the state vector always points towards the North or South pole of the sphere, i.e., along the $\hat{z}$ or $-\hat{z}$ direction.

When multiple qubits are brought together, their joint Hilbert space grows exponentially: a quantum computer with n qubits spans over 2^n basis states. Consequently, representing a general n -qubit state requires 2^n bits on a classical computer. This exponential scaling underlies the difficulty of classically simulating physical systems in nature, which are inherently quantum mechanical, and motivates the use of quantum devices for such tasks.

More importantly, quantum systems can exhibit *entanglement*, a unique quantum correlation phenomenon with no classical analogue. An entangled state of a bipartite system AB cannot be decomposed into independent states of individual subsystems A and B ; instead, information is encoded non-locally across the entire system. As a result, even complete knowledge of the composite state AB does not imply complete knowledge of its individual parts. This counter-intuitive property underpins many quantum algorithms and protocols, especially in quantum communication and cryptography.

These properties - quantum interference, superposition, and entanglement, are the key resources that distinguish quantum computation from classical computation and enable potential computational advantages. Nevertheless, it is important to emphasise that quantum computers do not universally outperform classical ones. For many problems, classical algorithms remain more efficient, and the overhead associated with QEC is substantial. Rather than replacing classical computation,

¹Up to a global phase, which has no physical significance, and thus can be neglected

quantum computing should be best viewed as a complementary paradigm, suited to specific classes of problems where quantum mechanical effects can be used as computational resources.

1.2 Physical realisation of a quantum computer

While the theoretical foundations of quantum computing are well established, its physical realisation has proven extremely challenging. A widely accepted set of guidelines for implementing quantum computation was proposed in 2000 by David DiVincenzo and known as the DiVincenzo criteria [10]. These requirements provide a useful framework for assessing suitable candidate physical platforms. In particular, a viable quantum computer should:

1. be scalable with well-defined qubits,
2. allow qubit initialisation into a known state,
3. have coherence times much longer than gate operation times,
4. support a universal set of quantum gates, and
5. enable high-fidelity, qubit-specific measurements.

Various physical platforms have been proposed as viable options that satisfy these requirements, including trapped ions [11, 12], neutral atoms [13], spin qubits [14], photonic systems [15], superconducting quantum circuits [16] and topological qubits [17]. However, so far, none of them can meet all of the criteria perfectly.

Trapped ions, for example, come as a natural implementation due to their intrinsically quantum nature. In such systems, individual ions are confined in electromagnetic traps with two long-lived states (ground state and excited state) and served as the qubits [12]. These systems exhibit exceptionally long coherence times (longer than a second for certain ions) and high gate fidelities. However, they are limited by relatively slow gate operations (typically in the microsecond regime) and challenges in scaling to large qubit numbers. Furthermore, it is difficult to perform two-qubit gates with trapped ions, since they require sufficiently strong interactions between two atoms [12].

Spin qubits in semiconductor quantum dots encode information in the spin degree of freedom of electrons or nuclei [14]. A major advantage of this approach is its compatibility with established semiconductor fabrication techniques, offering a potential path towards large-scale integration. Nevertheless, maintaining long coherence times and achieving precise control across many qubits remain significant challenges.

Superconducting circuits, pursued by big companies like Google and IBM, are arguably one of the most widely studied and mature platforms up to date. In these systems, information is encoded in the discrete energy levels of a nonlinear macroscopic electrical circuit, which acts like an anharmonic oscillator. This platform offers fast gate operations on the order of tens of nanoseconds and is highly compatible with modern microfabrication techniques from the semiconductor industry. While scaling a superconducting quantum processor itself is relatively simple, this is not the case for its operation-enabling infrastructure due to the complicated cryogenic setup and the large number of control and readout lines required. Moreover, superconducting qubits typically suffer from short coherence times, which are in the

order of tens to hundreds of microseconds.

Other platforms also offer promising directions. Neutral atoms, trapped in optical lattices or tweezers, provide excellent scalability and long coherence times [13], while photonic systems are naturally suitable for quantum communication and operate at room temperature [15]. Topological qubits, although still in an early stage, aim to encode information in intrinsically protected states, offering robustness against certain errors [17].

There are several common challenges across all platforms. Decoherence imposes a strict limit on computation time: qubits need to retain their quantum states long enough, relative to gate operation times, for quantum algorithms to be executed reliably.

Strongly related is the requirement of high gate fidelity (or low gate error), i.e., quantum gates must closely match their ideal cases. To allow for reliable FTQC, gate errors must fall below certain thresholds, such as those required by the surface code [18–20]. Achieving such precision, especially for two-qubit gates, is technically demanding, as control imperfections, noise, and calibration errors all contribute to gate infidelity.

As the system scales to a larger number of qubits, additional system-level challenges, such as unwanted interactions between qubits, may arise and degrade gate performance. These issues highlight the importance of robust control techniques and careful system design, and will be discussed in more detail in this thesis.

Given these persisting problems, no single quantum computing platform has yet demonstrated a clear advantage over the others. The remainder of this thesis focuses solely on superconducting qubits.

1.3 Superconducting quantum computing

The development of superconducting qubits dates back to the 1980s following the discovery of quantum phenomena in superconducting electrical circuits, such as circuit quantisation [21] and macroscopic quantum tunnelling [22]. These findings established the ability to treat superconducting circuits as *artificial atoms* with tunable properties such as transition frequencies. The first superconducting qubit - the Cooper pair box, was realised in 1999 by Nakamura *et al.* based on these ideas [23]. Variants operating in different parameter regimes soon followed, giving rise to charge [24], flux [25], and phase qubits [26], each exploiting different degrees of freedom of the superconducting circuit. A further advance came in the early 2000s with the development of circuit quantum electrodynamics (cQED) [27], in which superconducting qubits are coupled to microwave resonators, enabling high-fidelity quantum non-demolition (QND) readout. Another major breakthrough was the introduction of the transmon qubit by Koch *et al.* [28], which significantly improved coherence by suppressing charge noise sensitivity. The transmon, along with its cross-shaped planar variant Xmon [29], has since become the dominant superconducting qubit architecture, forming the basis of most contemporary quantum processors, including the one used in this thesis. In parallel, alternative designs with promising properties have also been explored, such as fluxonium qubits [30], bosonic modes [31], and qudits [32].

1.3.1 The Josephson junction (JJ)

At the heart of all superconducting circuits is the JJ. It comprises two superconducting electrodes separated by a thin barrier, which could be an insulator, a normal conductor, a constriction, or a grain boundary (Fig. 1.2). When the barrier is sufficiently thin, the two superconducting islands can couple weakly, meaning that their macroscopic wavefunctions can overlap, allowing Cooper pairs to tunnel through the barrier and sustain a dissipationless supercurrent across the junction.

The dynamics of a JJ is governed by the two Josephson equations. The first one, known as the DC Josephson effect, describes the supercurrent I_s flowing through the junction:

$$I_s = I_c \sin \phi, \quad (1.3)$$

with I_c the critical current of the junction and ϕ the superconducting phase difference across the junction². The second relation, known as the AC Josephson effect, says that a voltage V develops across the junction if the phase difference ϕ evolves in time:

$$\frac{d\phi}{dt} = \frac{2e}{\hbar} V. \quad (1.4)$$

From these two equations, the JJ appears as a nonlinear, dissipationless inductor with an effective inductance L_J given by

$$L_J = \frac{\hbar}{2eI_c \cos \phi} = \frac{L_{J0}}{\cos \phi}. \quad (1.5)$$

Unlike a linear inductance, which is current-independent, the Josephson inductance L_J depends on the phase difference, and hence the current (through Eq. 1.3).

The energy stored in the JJ is calculated by integrating the product of voltage and current with respect to time, $E = \int_{-\infty}^t V(t') I(t') dt'$. Using the two Josephson relations yields

$$E = \frac{\hbar I_c}{2e} (1 - \cos \phi) = E_J (1 - \cos \phi), \quad (1.6)$$

with $E_J = \hbar I_c / 2e$ being the Josephson energy. The resulting cosine potential (Fig. 1.3) is fundamentally different from the quadratic potential of a harmonic oscillator (e.g., an LC circuit) and leads to an anharmonic energy spectrum. When incorporated into a circuit, this anharmonicity breaks the equally spaced energy level structure, allowing separation of the two lowest levels that define the computational basis $\{|0\rangle, |1\rangle\}$ from higher ones.

An important extension of the JJ is the superconducting quantum interference device (SQUID), which consists of two JJs arranged in a loop. A symmetric SQUID (i.e., one with identical JJs) behaves like a single JJ with a tunable critical current I_c^{SQUID} [33]:

$$I = I_c^{\text{SQUID}} \sin \phi, \quad (1.7)$$

$$I_c^{\text{SQUID}} = 2I_c \cos \left(\frac{\pi \Phi}{\Phi_0} \right), \quad (1.8)$$

²Here the contribution from the vector potential $\mathbf{A}$ is neglected.

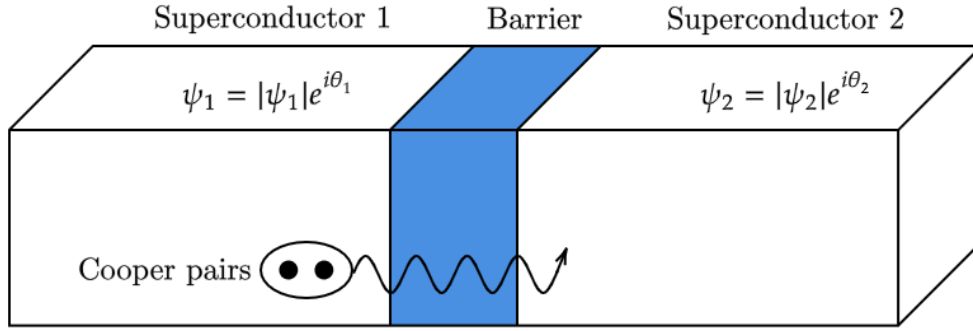


Figure 1.2: Illustration of a JJ. The two superconducting electrodes with macroscopic wavefunctions $\psi_n = |\psi_n|e^{i\theta_n}$ ($n = 1, 2$) are separated by a thin barrier (blue) that allows Cooper pairs to tunnel through. Adapted from [34].

where I_c is the critical current of each individual JJ making up the SQUID, Φ is the externally applied magnetic flux through the loop, and $\Phi_0 = h/2e$ is the flux quantum. This flux tunability provides a powerful control knob over different circuit parameters such as qubit frequencies and coupling strength. The SQUID is therefore widely used in many superconducting quantum processor architectures.

1.3.2 From Cooper pair box to transmon qubit

The simplest and earliest implementation of a superconducting qubit is the Cooper pair box (or charge qubit) [23], which can be understood as an LC circuit where the linear inductor is replaced by a JJ (Fig. 1.4a). The presence of the nonlinear JJ introduces anharmonicity to the energy potential (Fig. 1.3). To control the charge on the superconducting island, the qubit is capacitively coupled to a gate voltage V_g , which induces an offset charge $n_g = -C_g V_g / 2e$ on the island. By tuning this offset charge, the energy levels of the system can be adjusted, as shown in Fig. 1.5. The Hamiltonian in this case becomes

$$\hat{H} = 4E_C(\hat{n} - n_g)^2 - E_J \cos \hat{\phi}, \quad (1.9)$$

where E_C is the charging energy of an electron:

$$E_C = \frac{e^2}{2C} = \frac{e^2}{2(C_J + C_g)}. \quad (1.10)$$

While this dependence of energy levels on n_g is useful, it also makes the Cooper pair box extremely sensitive to charge noise, as evident from the curvature of the plot in Fig. 1.5 for the case of $E_J = E_C$. Fluctuations in the qubit energy levels will lead to rapid dephasing, and thus severely limit coherence times.

One solution to this problem is the transmon qubit, which is formed by shunting the JJ with a large capacitance C_S (Fig. 1.4b). The added capacitance reduces the charging energy E_C , thereby pushing the system into the regime $E_J \gg E_C$. In this limit, the qubit's energy levels become independent of n_g (Fig. 1.5), significantly suppressing its susceptibility to charge noise. The first two transition energies E_{01}

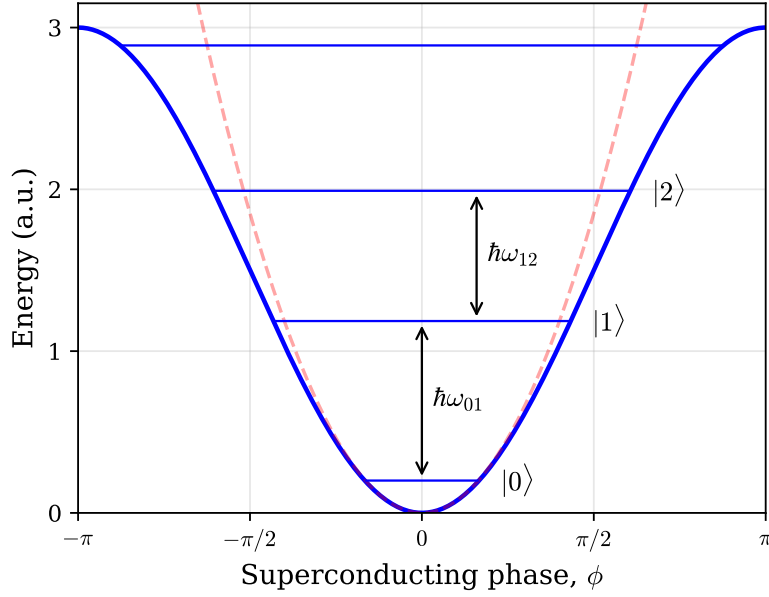


Figure 1.3: The anharmonic energy potential of the transmon qubit. The nonlinear Josephson inductance in a superconducting circuit results in a sinusoidal energy potential (solid blue) instead of parabolic as in a harmonic oscillator (dashed red), leading to non-equidistant energy levels. It is thus possible to form a computational basis from the two lowest eigenstates. Adapted from [35].

and E_{12} of a transmon qubit are approximately given by

$$E_{01} = \hbar\omega_{01} = \sqrt{8E_J E_C} - E_C, \quad (1.11)$$

$$E_{12} = \hbar\omega_{12} = E_{01} - E_C = \sqrt{8E_J E_C} - 2E_C, \quad (1.12)$$

where ω_{01} and ω_{12} are the first and second transition frequency of the qubit, respectively. In a transmon structure with only one JJ (Fig. 1.4b), these frequencies are fixed. However, if the JJ is replaced by a symmetric SQUID (Fig. 1.4c), the frequencies become tunable via an external magnetic field applied through the loop. This magnetic flux will modify the critical current of the SQUID (according to Eq. 1.8), which in turn will change the Josephson energy (Eq. 1.6), and hence the transition frequencies (Eqs. 1.11, 1.12).

The anharmonicity of a qubit, α is defined as the difference between its first two transition frequencies:

$$\alpha = \omega_{12} - \omega_{01} = -\frac{E_C}{\hbar} \quad (1.13)$$

This expression shows the trade-off in a transmon qubit: increasing the E_J/E_C ratio suppresses charge noise sensitivity but at the same time, reduces anharmonicity. In the extreme of very high E_J/E_C ratios, the transition frequencies become almost identical and the system approaches a harmonic oscillator. Leakage, i.e., the loss of quantum information from the computational subspace into higher excited states, then becomes a serious problem.

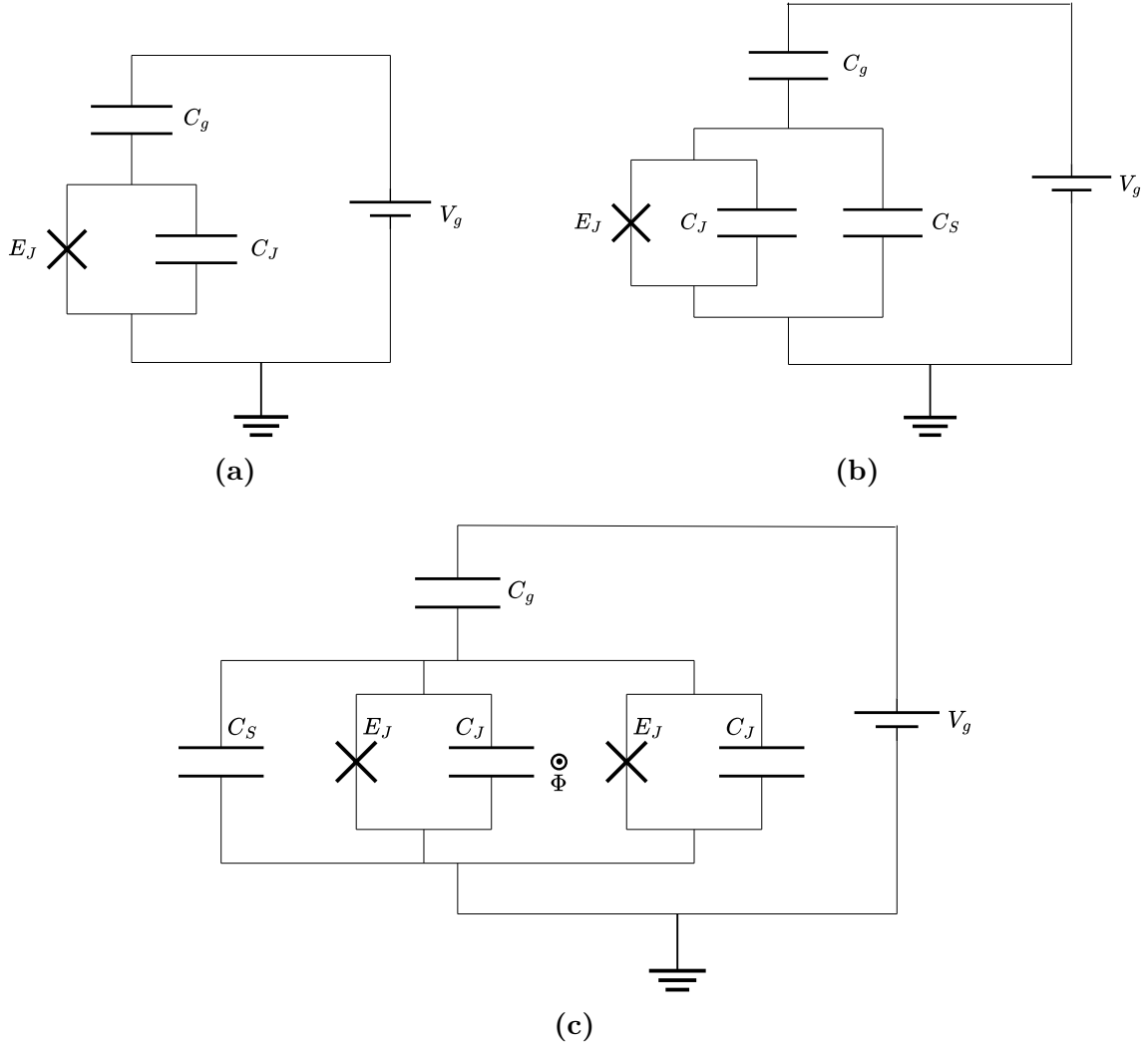


Figure 1.4: Circuit diagrams of superconducting qubits. (a) Charge qubit: The JJ (denoted by a cross) has a self-capacitance C_J and Josephson energy E_J . The upper node is coupled to a voltage source V_g via a gate capacitor C_g . (b) Fixed-frequency transmon qubit: The JJ is shunted with a large capacitor C_S to increase the E_J/E_C ratio. The transition frequencies are fixed in this case. (c) Flux-tunable transmon qubit: The JJ is replaced with a SQUID so that the qubit frequencies can be adjusted by applying an external magnetic flux Φ .

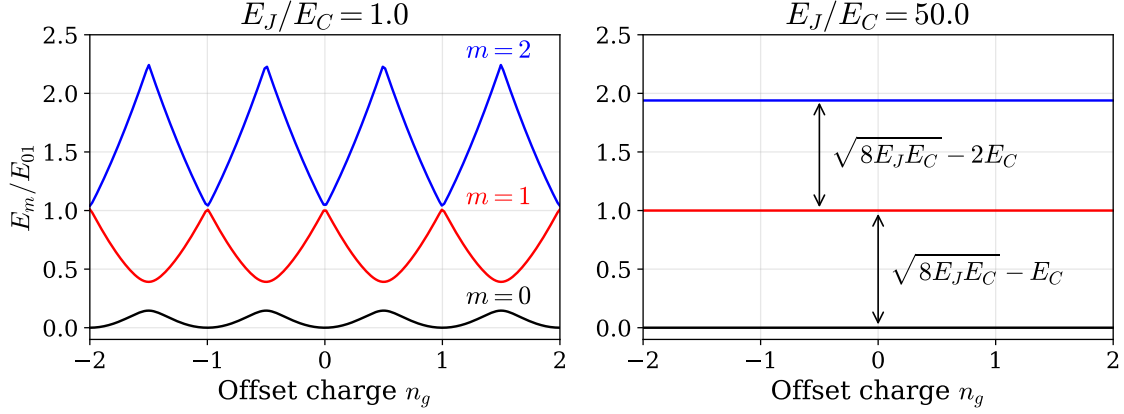


Figure 1.5: The dependence of energy levels on the offset charge. The energy of the first three levels, E_m ($m = 0, 1, 2$) in units of the first transition energy, E_{01} is plotted as a function of the offset charge n_g for $E_J = E_C$ (charge qubit) and $E_J = 50E_C$ (transmon qubit). Adapted from [28].

1.3.3 Gate operations in superconducting circuits

Single-qubit gates operate on individual qubits and correspond to rotations of the state vector on the Bloch sphere. In superconducting circuits, they are realised by applying microwave pulses on resonance, i.e., at the qubit's resonance frequency, ω_{01} . The driven qubit can be described by the Hamiltonian

$$\hat{H} = -\frac{\hbar\omega_{01}}{2}\hat{\sigma}_z - \Omega(t)\cos(\omega_d t + \phi)\hat{\sigma}_x, \quad (1.14)$$

where $\Omega(t)$ is the pulse envelope, $\omega_d \approx \omega_{01}$ is the drive frequency, and ϕ is the pulse phase. At resonance ($\omega_d = \omega_{01}$, i.e., detuning $\Delta = \omega_d - \omega_{01} = 0$), and in a frame rotating at ω_d , the Hamiltonian becomes

$$\hat{H}_{\text{rot}} = \frac{\Omega(t)}{2}(\hat{\sigma}_x \cos \phi + \hat{\sigma}_y \sin \phi). \quad (1.15)$$

This expression makes it clear that arbitrary rotations about any axes in the xy-plane can be implemented by controlling the duration, amplitude, and phase of the pulse. Rotations about the z-axis are typically realised by virtual Z gates, i.e., by updating the reference phase of subsequent pulses without applying any additional physical pulses [16]. As mentioned in Section 1.2, achieving a low gate error (or equivalently, a high gate fidelity) is essential for reliable quantum computation. In practice, this requires careful calibration of pulse parameters to mitigate imperfections such as leakage.

Two-qubit gates are required to generate entanglement between qubits and complete a universal gate set. In superconducting circuits, such gates are typically implemented via capacitive or inductive coupling, either directly or mediated by a coupler [16]. A widely used approach in fixed-frequency transmon systems is the cross-resonance gate, in which a microwave drive is applied to one qubit at the resonance frequency of a neighbouring qubit, inducing an effective conditional interaction [36]. Other common implementations include controlled-phase (CZ) and

iSWAP gates, which rely on tuning the qubits into resonance or activating specific coupling pathways. Although many types of entangling two-qubit gates exist, any of them is sufficient to complete the universal quantum gate set, as all others can be decomposed into combinations of that gate and arbitrary single-qubit gates. This thesis focuses primarily on single-qubit gate operations.

1.3.4 Dispersive readout of superconducting qubits

Superconducting qubits are most commonly measured using the so-called dispersive readout within the framework of cQED [27]. In this approach, the qubit is coupled to a microwave cavity (or resonator). The dynamics of the combined qubit-cavity system is described by the Jaynes-Cummings Hamiltonian:

$$\hat{H} = -\frac{\hbar\omega_{01}}{2} + \hbar\omega_c\hat{a}^\dagger\hat{a} + \hbar g(\hat{\sigma}_+\hat{a} + \hat{\sigma}_-\hat{a}^\dagger), \quad (1.16)$$

where ω_c denotes the resonator frequency, g is the qubit-cavity coupling strength, $\hat{a}^\dagger$ ($\hat{a}$) is the bosonic creation (annihilation) operator, and $\hat{\sigma}_+$ ($\hat{\sigma}_-$) is the raising (lowering) operator. In the so-called dispersive regime where the detuning between the qubit and resonator frequencies, $\Delta = \omega_{01} - \omega_c$ is much larger than the coupling strength ($|\Delta| \gg g$), the interaction can be approximated by an effective Hamiltonian of the form

$$\hat{H}_{\text{disp}} = -\frac{\hbar}{2}(\omega_{01} + \chi)\hat{\sigma}_z + \hbar(\omega_c - \chi\hat{\sigma}_z)\hat{a}^\dagger\hat{a}, \quad (1.17)$$

with $\chi = g^2/\Delta$ being the dispersive shift. The second term of this Hamiltonian implies that the resonator frequency is shifted by an amount $\pm\chi$ depending on whether the qubit is in the ground state $|0\rangle$ or the excited state $|1\rangle$ (Fig. 1.6a). As a result, by probing the resonator with a microwave signal and measuring its transmission or reflection response, one can infer the state of the qubit without having to actually measure it. This is a manifestation of QND readout, as it does not induce any transitions between the qubit states.

Dispersive readout enables two-state single-shot readout (SSRO), in which the two states $|0\rangle$ and $|1\rangle$ can be distinguished in a single run. Each run produces a complex signal that can be represented as a coordinate in the IQ plane (Fig. 1.6b). Repeating the experiment many times gives two distinct clusters of points corresponding to the two states. If the clusters are well separated, each measurement outcome can then be assigned to its correct state with high fidelity. This approach can also be extended to three-state SSRO by including the second excited state $|2\rangle$, which is relevant for detecting leakage in a weakly anharmonic system like the transmon. More details on SSRO, as well as its use in randomised benchmarking (RB) for characterising gate performance and leakage errors, are given in Appendix A.

1.3.5 Qubit decoherence

As briefly mentioned in Section 1.2, decoherence is one of the main factors limiting the performance of quantum computers. In practice, a qubit cannot be perfectly isolated from its surrounding environment, as some coupling is necessary for control

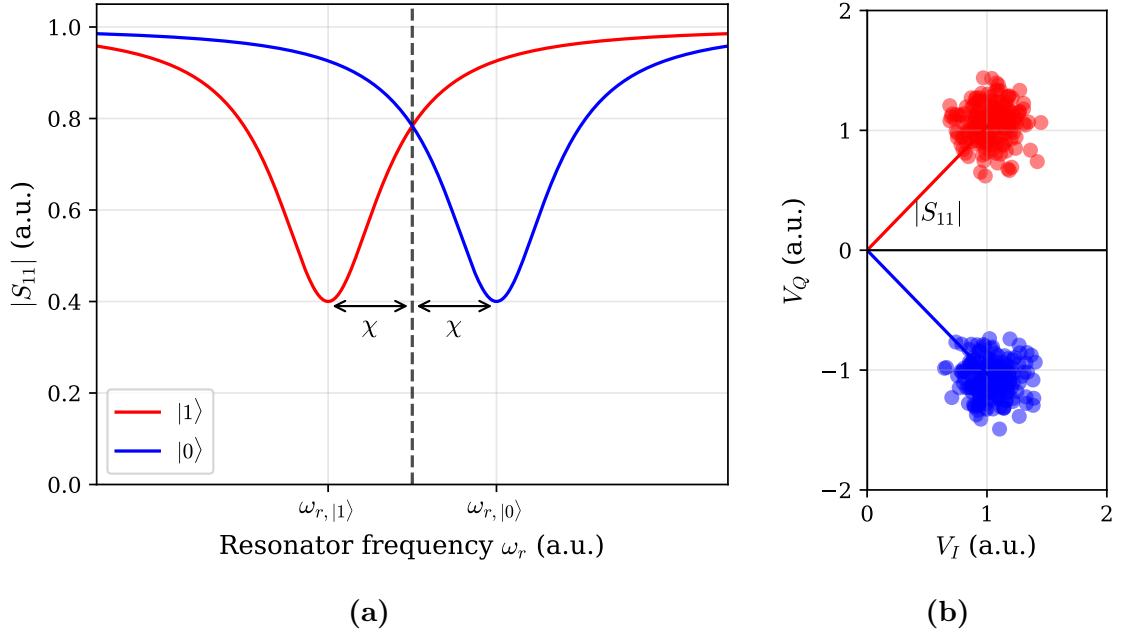


Figure 1.6: Dispersive readout of superconducting qubits. (a) The magnitude of the reflection response of the resonator, $|S_{11}|$ changes depending on the state of the qubit. The shift corresponds to an angular frequency of 2χ . (b) Complex plane representation, with the two axes being the real and quadrature component of S_{11} . In this case, the two states can be best distinguished if the resonator is probed in between them (the dashed line in (a)). Adapted from [35].

and measurement. This inevitable interaction introduces noise and leads to the loss of quantum information.

One way to model decoherence is to treat the qubit as interacting with a reservoir (or thermal bath). Longitudinal noise³ leads to pure dephasing characterised by a timescale T_2^* (or T_ϕ). In this case, the diagonal elements of the density matrix remain unchanged, while the off-diagonal elements decay exponentially at a rate $\Gamma_2^* = 1/T_2^*$. Physically, this corresponds to the loss of phase coherence without changing the energy level populations. Dephasing is typically dominated by low-frequency noise sources such as charge noise, flux noise, and critical current noise [37].

On the other hand, transverse noise³ induces transitions between energy levels, leading to relaxation and excitation at rates $\Gamma_\downarrow$ and $\Gamma_\uparrow$, respectively. The total mixing rate is defined as $\Gamma_1 = \Gamma_\downarrow + \Gamma_\uparrow$. In superconducting qubits operating at mK temperatures (where $k_B T \ll \hbar\omega_{01}$), excitation is strongly suppressed and relaxation dominates. Thus, $T_1 = 1/\Gamma_1$ is referred to as the relaxation time of the qubit. This process is often limited by quasi-particle tunneling [38–40], dielectric loss in the material [41, 42], radiation [43], Purcell decay [44, 45], and magnetic field noise [46]. Importantly, transverse noise also contributes to dephasing with rate $\Gamma_1/2$. Combining the contributions from both transverse and longitudinal noises, the total

³Here, *longitudinal* and *transverse* refer to the direction of the coupling between the qubit and the reservoir.

dephasing time T_2 is then given by

$$\frac{1}{T_2} = \frac{1}{2T_1} + \frac{1}{T_2^*}. \quad (1.18)$$

Mitigating decoherence requires careful circuit design, improved materials and fabrication processes, as well as effective shielding from radiative and magnetic field sources [47]. The experiments used to measure T_1 and T_2 are discussed in detail in Section A.3 of Appendix A.

1.4 Scaling superconducting quantum processors

1.4.1 Motivation for scaling and 3D integration

Many quantum algorithms of practical interest require deep circuits, while the implementation of QEC imposes a significant overhead in the number of physical qubits, potentially requiring up to millions of them for large-scale FTQC. As a result, the performance of current quantum devices is no longer limited solely by single-qubit coherence or gate fidelity, but increasingly by the ability to scale them while preserving these metrics.

However, this is not a trivial problem. As the system grows, the routing of the microwave circuitry becomes increasingly complex. In conventional planar superconducting architectures, wiring crowding quickly emerges, as signal lines must be carefully routed to avoid crossings and unintended interactions. Moreover, the increased wiring density exacerbates signal crosstalk between control lines and qubits [47] (see Section 1.5).

Various techniques have been employed to mitigate these issues, such as the use of airbridges and other crossover structures [47–49], as well as enclosing transmission lines within grounded “tunnels” to improve electromagnetic shielding [48–50]. While effective for small-scale systems, these approaches become increasingly difficult to implement as the number of qubits grows [47, 48]. These limitations motivate the exploration of alternative architectures that leverage the third dimension of space, such as flip-chip architectures.

1.4.2 Flip-chip integration

Originally developed for semiconductor microelectronics, flip-chip architectures have been adapted to superconducting quantum processors and now stand out as one of the most promising multi-chip approaches for addressing scaling problems [48, 50–53]. In this technology, the system is typically divided into two separate chips: a quantum chip (Q-chip), which hosts the qubits and couplers, and a control chip (C-chip) containing the microwave circuitry. These two chips are then electrically interconnected via superconducting bump bonds. An example of a two-qubit flip-chip device fabricated in collaboration between the QC2 group at Chalmers University of Technology and VTT is shown in Fig. 1.7 [48].

This separation of functionality provides several advantages. By relocating the microwave circuitry to a dedicated chip, the spacing between signal lines can be

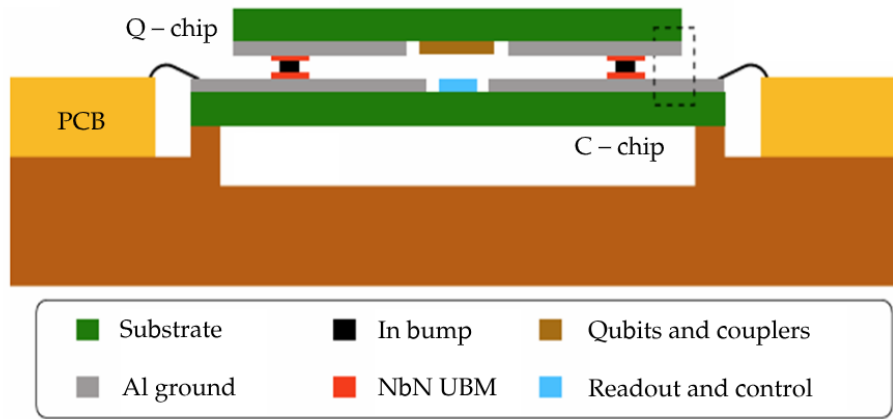


Figure 1.7: Illustration of a two-qubit flip-chip superconducting quantum processor. The Q-chip contains the qubits and couplers, while the C-chip accommodates the readout resonators, feedthrough transmission lines, and control lines. The two chips are interconnected through indium (In) bumps, with an under-bump metallisation (UBM) layer separating each bump from the aluminum (Al) ground planes. The control lines and ground plane on the C-chip are wire-bonded to a printed circuit board (PCB), which is mounted together with the module inside a copper (Cu) sample box. Adapted from [48].

increased, thereby significantly reducing on-chip crowding and thus suppressing crosstalk. Additionally, the modular nature of flip-chip integration allows independent optimisation and fabrication of the quantum and control layers, reducing the risk of degrading qubit performance [48]. These benefits, however, come with extra complexity in chip design, fabrication and thermal management.

Nevertheless, flip-chip architectures, and more broadly 3D integration, provide a promising pathway towards scalable superconducting quantum processors. The experiments presented in this thesis were performed on a 25-qubit flip-chip processor fabricated at Chalmers, which is described in more detail in Section 3.1.

1.4.3 Frequency-multiplexed qubit control and readout

In conventional architectures, each qubit is typically connected to its own dedicated microwave control and readout lines, leading to a rapid increase in hardware complexity as the number of qubits grows. Multiplexing alleviates this bottleneck by allowing multiple qubits to share common control and readout infrastructure.

In the case of multiplexed readout, a shared transmission line carries a multi-tone signal, where each frequency component addresses a specific qubit via its corresponding readout resonator. Since each resonator is designed to have a distinct frequency, the response of individual qubits can be distinguished in the frequency domain. This approach significantly reduces the number of required feedlines and associated hardware. However, maintaining high readout fidelity requires careful engineering of resonator frequencies and filtering to minimise spectral overlap and interference.

A similar concept can be applied to frequency-multiplexed qubit control. In this scheme, multiple qubits are driven through a shared microwave control line using

pulses at different frequencies, each tuned to the transition frequency of a target qubit. While this reduces the number of control lines, it can introduce unwanted transitions in non-target qubits if the frequency detunings between qubits are not sufficiently larger than the bandwidth of the microwave pulse. This leads to signal crosstalk, and thus to an increase in gate errors, which becomes more pronounced as the system size increases.

The 25-qubit quantum processor used in this work employs multiplexed readout, where every five qubits share a common transmission line. Each qubit, however, still has its own control lines. More details are given in Section 3.1.

1.5 Crosstalk in a superconducting quantum processor

In general, crosstalk refers to unintended interactions among qubits, or between qubits and their control and readout circuitry. Experimentally, it manifests as a degradation in performance when gate operations or measurements are executed simultaneously on multiple qubits compared to when they are done individually. This effect becomes increasingly important as quantum processors scale up and require a growing number of control and readout lines [54]. In conventional architectures, these lines are often routed from the chip periphery via wire bonds to the external circuitry [50]. As the routing density increases, maintaining sufficient isolation between the lines and suppressing unwanted signal couplings become increasingly challenging [47, 48, 54].

Crosstalk is one of the leading sources of coherent gate errors in superconducting quantum circuits [55]. This is particularly critical in the context of QEC, where error rates must be kept below certain thresholds [18–20, 56]. Moreover, the threshold theorem typically assumes that the errors are local and weakly correlated (e.g., bit-flip and phase-flip errors) [57, 58]. Crosstalk can introduce global and correlated errors that violate these assumptions, potentially affecting the successful implementation of QEC [19, 55–57].

Crosstalk can be broadly categorised into two main types: classical crosstalk (or signal crosstalk), which originates from unwanted classical electromagnetic interactions between the drive signals of neighbouring qubits, and quantum crosstalk, which arises from spurious residual Hamiltonian couplings between qubits [56–58]. Other forms of crosstalk, such as measurement crosstalk, where the readout of one qubit inadvertently influences the outcome of another, can also contribute to performance degradation [59, 60].

1.5.1 Classical crosstalk

There are two types of classical crosstalk: XY (microwave, or qubit-drive) crosstalk and Z (or flux) crosstalk, each corresponding to a type of control line in the processor. XY crosstalk occurs between microwave drive (XY) lines used for qubit control, while Z crosstalk occurs between flux-bias (Z) lines used to tune the qubit or coupler frequencies. This thesis focuses solely on XY crosstalk.

1.5.1.1 XY crosstalk

XY crosstalk occurs when a microwave drive pulse transmitted through a certain control line to its corresponding qubit unintentionally reaches a neighbouring qubit and induces spurious transitions [57]. An off-resonant crosstalk drive can also induce a shift in the frequency of the affected qubit, resulting in gate errors [54]. This is known as the AC Stark effect.

As shown in Fig. 1.8, XY crosstalk can arise from several mechanisms, including stray coupling between drive lines [61], direct capacitive coupling between qubits and control lines [50], and indirect coupling mediated by packaging modes or parasitic circuit elements (e.g., airbridges, metal enclosure, PCB, and wire bonds) [61–63]. In densely integrated systems, these effects are more pronounced due to the proximity of the components and the complexity of the electromagnetic environment.

The impact of XY crosstalk is evident in hardware characterisation protocols like RB [56, 58] as an increase in gate errors when multiple qubits are driven simultaneously compared to when they are driven one at a time. This is particularly relevant for two qubits whose transition frequencies are close to one another - a phenomenon referred to as frequency collision [57]. If two qubits are far detuned from one another, even when an unwanted microwave signal is present due to limited physical isolation, no transition would be induced, and the AC Stark effect would also be suppressed by the large detuning.

In practice, however, especially in a large-scale quantum processor with many qubits, the opposite is often true. Frequency collisions are difficult to avoid since transmons operate within a relatively narrow frequency window, typically from 4 GHz to 8 GHz [54], and the frequencies of two neighbouring qubits need to be sufficiently close to each other to enable two-qubit gates (see Section 3.2). These collisions could occur either between the first transition frequencies of two qubits ($f_{01} - f_{01}$), or between the first transition frequency of one qubit and the second transition frequency of another ($f_{01} - f_{12}$).

In the first case, the resulting errors remain within the computational subspace and, therefore, can be corrected through unitary transformations [56]. In contrast, the second type of collision leads to leakage errors, which are particularly problematic because they cannot be corrected by unitary transformations in the computational subspace and require substantial additional overhead in fault-tolerant protocols [56, 57, 64–66]. Moreover, even when the qubit relaxes back to its computational basis (i.e., seepage), the quantum information lost during leakage is not recovered [64, 66].

Since crosstalk dynamics are based on qubit pairs, in this thesis, the qubit that is intentionally driven is referred to as the *control qubit*, while the one being inadvertently affected is the *target qubit*. XY crosstalk between a control qubit i and a target qubit j can then be described by a complex crosstalk coefficient β_{ji} (see Section 2.1):

$$\beta_{ji} = |\beta_{ji}|e^{i\varphi_{ji}}. \quad (1.19)$$

The amplitude and phase of β_{ij} can be experimentally measured using Rabi-like [20, 50, 55, 67, 68], Ramsey-like [67], or spin-echo sequences [54], as discussed further in Section 2.2 and Chapter 4. Current superconducting devices typically exhibit XY crosstalk levels above 1 % [68].

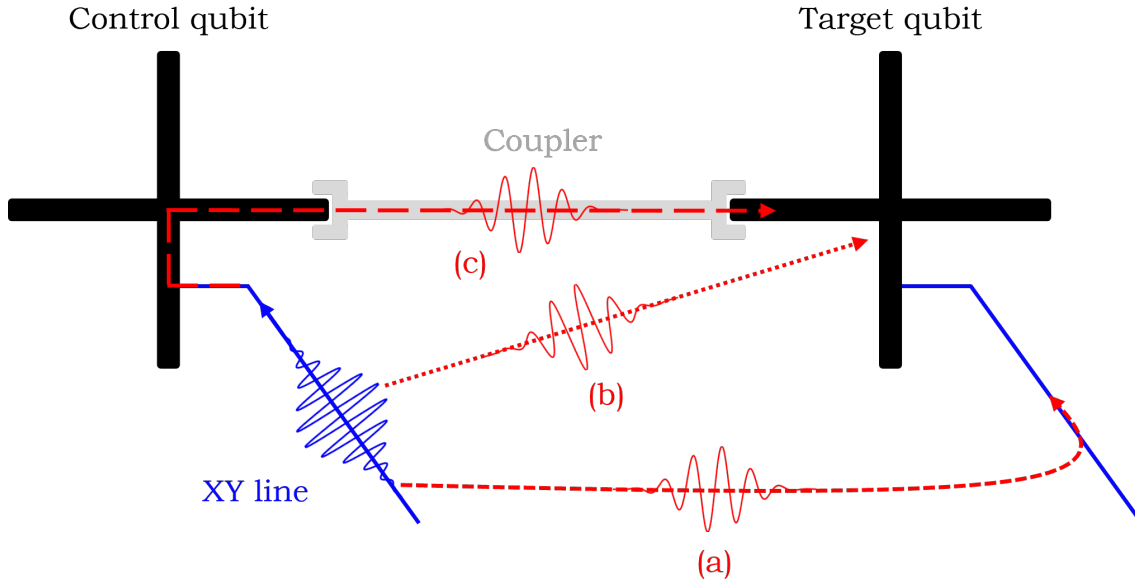


Figure 1.8: Possible XY crosstalk channels between two superconducting qubits. XY crosstalk in a superconducting circuit can arise from (a) stray coupling between two XY lines (blue), (b) direct capacitive coupling between an XY line and a nearby non-target qubit (black), or (c) indirect coupling via packaging modes or circuit elements, such as couplers (grey). Adapted from [61].

Mitigation strategies for XY crosstalk include both hardware- and software-level approaches. On the hardware side, improved qubit circuit design, increased spacing, and shielding between control lines can reduce unwanted coupling [50, 61]. On the software side, a widely used method is active compensation, where an additional microwave pulse with the same frequency and amplitude but opposite phase is applied to cancel out the crosstalk signal [54, 67, 68]. Other approaches include pulse shaping techniques that suppress transitions at certain frequencies [69, 70], and optimisation of parametrised single qubit gates [56, 71]. In addition to these mitigation strategies, frequency allocation in the quantum processor is a crucial design consideration, as the frequency separation between neighbouring qubits should be maximised to avoid both $f_{01} - f_{01}$ and $f_{01} - f_{12}$ collisions [47, 50, 70]. This thesis focuses on the active compensation approach, which will be discussed in more detail in Sections 2.3 and 4.3.

1.5.2 Quantum crosstalk

Unlike classical crosstalk, quantum crosstalk can be present even in the absence of frequency collisions. It arises from always-on parasitic ZZ interactions between qubits, which shift the transition frequency of one qubit depending on the state of its neighbour, producing correlated, state-dependent phase errors. Quantum crosstalk will not be discussed in detail in this thesis.

1.6 In this Master's thesis

1.6.1 Scope & objectives

This thesis addresses XY crosstalk, a central limitation to high-fidelity simultaneous gate operations in large-scale superconducting quantum processors. In particular, the objectives of the thesis are to:

- develop experimental protocols for measuring both the amplitude and phase of the complex XY crosstalk coefficient for $f_{01} - f_{01}$ and $f_{01} - f_{12}$ crosstalk;
- integrate these measurement protocols into the existing automated calibration software framework;
- implement an active signal-level crosstalk compensation scheme based on the measured crosstalk coefficients; and
- evaluate the effectiveness of the compensation scheme using isolated and simultaneous RB, with the average single-qubit gate error and leakage rate as performance metrics.

The scope of the thesis is limited to XY crosstalk. Other forms of crosstalk, such as Z crosstalk, quantum crosstalk, and measurement crosstalk, are outside the scope of this work. Likewise, the thesis focuses on calibration and control techniques rather than hardware re-design or packaging improvements. All experiments were performed on Chalmers' 25-qubit fixed-frequency transmon flip-chip quantum processor, which had already been fabricated, packaged, and installed in a dilution refrigerator prior to the start of the project. Therefore, no fabrication or cryogenic setup was involved.

This work aims to provide an automated workflow for characterising and mitigating XY crosstalk within the software framework of the QC2 group and to improve simultaneous single-qubit gate performance on the 25-qubit chip. The thesis also aims to contribute to the general quantum community by presenting a simple and efficient measurement method adapted to fixed-frequency transmon qubits and by establishing a consistent and clear framework for distinguishing the different quantities associated with XY crosstalk, such as the different crosstalk coefficients.

1.6.2 Thesis outline

The thesis is organised as follows:

- **Chapter 2** provides a theoretical framework for describing XY crosstalk and its compensation. A brief literature review on existing XY crosstalk measurement methods is also included in this chapter.
- **Chapter 3** shows the architecture and frequency allocation scheme of the 25-qubit flip-chip superconducting quantum processor developed at Chalmers, as well as the cryogenic measurement setup and the software infrastructure used in this work.
- **Chapter 4** is the most important chapter of the thesis, where the crosstalk measurement and compensation protocols are described and discussed. Representative measurement and performance validation results are also shown for a selected qubit pair.
- **Chapter 5** summarises the main findings of the thesis and discusses directions

for future work.

- **Appendix A** describes the single-qubit gate calibration and characterisation workflow implemented in the group, which is highly relevant for the successful implementation of the crosstalk measurement and compensation protocols.
- **Appendix B** includes supplementary amplitude and phase measurement results, RB plots, and average single-qubit gate error matrices for other qubit pairs.

2

XY Crosstalk: Theory & Literature Review

This chapter introduces the theoretical framework used throughout this thesis to describe, measure, and mitigate XY crosstalk in superconducting quantum processors. Section 2.1 begins with a mathematical description of XY crosstalk, introducing the crosstalk matrix formalism and distinguishing between the bare and the effective crosstalk coefficient. Section 2.2 then reviews some existing experimental methods for measuring XY crosstalk, including Rabi-, Ramsey-, and spin-echo-based approaches. Finally, Section 2.3 discusses signal-level crosstalk compensation using the measured crosstalk matrix.

2.1 Mathematical formulation of XY crosstalk

2.1.1 Crosstalk coefficient & crosstalk matrix

Without loss of generality, first consider XY crosstalk between two qubits i and j . In an ideal scenario without any crosstalk, each qubit is driven only by the microwave signal applied through its dedicated XY line at its own transition frequency [68]:

$$\Omega'_i(t) = \Omega_{ii} e^{i(\omega_i t + \varphi_{ii})}, \quad (2.1)$$

$$\Omega'_j(t) = \Omega_{jj} e^{i(\omega_j t + \varphi_{jj})}, \quad (2.2)$$

where Ω , ω , and φ are the drive strength, angular frequency, and phase of the signal, respectively. The first index of the double subscript denotes the qubit receiving the signal, while the second index denotes the drive line from which the signal is sent. For example, Ω_{ii} represents the intended drive amplitude from line i to qubit i .

In the presence of crosstalk, however, each qubit experiences not only its intended drive but also an unwanted crosstalk contribution from the other qubit (Fig 2.1). The total microwave drives acting on the two qubits can thus be written as [68]:

$$\Omega_i(t) = \Omega_{ii} e^{i(\omega_i t + \varphi_{ii})} + \Omega_{ij} e^{i(\omega_j t + \varphi_{jj} + \varphi_{ij})}, \quad (2.3)$$

$$\Omega_j(t) = \Omega_{jj} e^{i(\omega_j t + \varphi_{jj})} + \Omega_{ji} e^{i(\omega_i t + \varphi_{ii} + \varphi_{ji})}, \quad (2.4)$$

where the notation ij (ji) is used for the signal leaking from drive line j (i) to qubit i (j). The crosstalk signal shares the frequency of its source but typically has a lower amplitude and an additional phase shift.

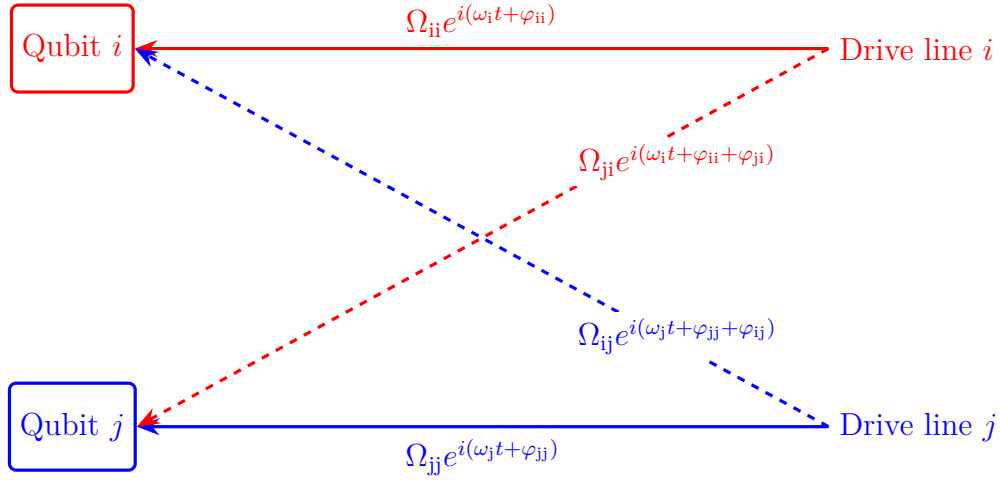


Figure 2.1: XY crosstalk between two qubits. The schematic shows two qubits i and j and their dedicated XY lines. Each qubit receives a drive signal at its resonance frequency from its own XY line (solid lines) and a crosstalk signal at the other qubit's frequency from the other line (dashed lines). Adapted from [68].

A complex crosstalk coefficient β_{ji} is then conveniently used to describe the microwave leakage from drive line i to qubit j :

$$\beta_{ji} = |\beta_{ji}|e^{i\varphi_{ji}}, \quad (2.5)$$

where the crosstalk magnitude $|\beta_{ji}|$ is defined as the ratio of the drive strength of the leaked signal from control line i onto qubit j to the intended drive strength from the same control line onto qubit i :

$$|\beta_{ji}| = \frac{\Omega_{ji}}{\Omega_{ii}}. \quad (2.6)$$

The phase φ_{ji} determines the relative phase offset introduced by the coupling between the two qubits and drive lines. Similarly, the crosstalk coefficient for leakage from line j to qubit i :

$$\beta_{ij} = |\beta_{ij}|e^{i\varphi_{ij}} = \frac{\Omega_{ij}}{\Omega_{jj}}e^{i\varphi_{ij}}. \quad (2.7)$$

Alternatively, the crosstalk level can be expressed in decibels as:

$$\Lambda = 10 \log_{10} |\beta|^2. \quad (2.8)$$

This figure of merit is often used in works that, e.g., compare crosstalk between different qubit pairs or investigate its dependence on qubit separation. In this thesis, however, the complex crosstalk coefficient β is more useful since it directly provides information on the compensation signal (see Section 2.3).

In matrix form, the intended drive signals, $\Omega'(t)$ and the actual signals received by the qubits, $\Omega(t)$ can be written as two column vectors [68]:

$$\Omega'(t) = (\Omega'_i(t) \quad \Omega'_j(t))^T, \quad (2.9)$$

$$\Omega(t) = (\Omega_i(t) \quad \Omega_j(t))^T. \quad (2.10)$$

The two are related by a forward crosstalk matrix C :

$$\Omega(t) = C\Omega'(t), \quad (2.11)$$

where

$$C = \begin{pmatrix} 1 & \beta_{ij} \\ \beta_{ji} & 1 \end{pmatrix}. \quad (2.12)$$

The same formalism extends naturally to a processor with N qubits. In that case, the total microwave drive experienced by qubit i is given by [68]:

$$\Omega_i(t) = \Omega'_i(t) + \sum_{j \neq i} \Omega_{ij} e^{i(\omega_j t + \varphi_{ji} + \varphi_{ij})}, \quad (2.13)$$

where the summation accounts for the crosstalk contributions from all other XY lines. The crosstalk matrix then becomes an $N \times N$ matrix whose diagonal elements are unity and off-diagonal elements correspond to the complex crosstalk coefficients between all qubit pairs:

$$C = \begin{pmatrix} 1 & \beta_{12} & \beta_{13} & \cdots & \beta_{1N} \\ \beta_{21} & 1 & \beta_{23} & \cdots & \beta_{2N} \\ \beta_{31} & \beta_{32} & 1 & \cdots & \beta_{3N} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \beta_{N1} & \beta_{N2} & \beta_{N3} & \cdots & 1 \end{pmatrix}. \quad (2.14)$$

This matrix completely characterises XY crosstalk present in the processor and forms the basis for the compensation method discussed in Section 2.3.

2.1.2 The effect of frequency detuning & effective crosstalk coefficient

The crosstalk coefficient β introduced in Section 2.1.1 so far only describes microwave leakage arising from the electromagnetic environment in the device. It depends on factors such as stray couplings, capacitive/inductive couplings, wiring, and packaging modes, but is independent of the qubits' frequencies themselves. It is thus referred to as the *bare* crosstalk coefficient.

During real gate operations, however, the leaked tone from control qubit i arriving at target qubit j is generally detuned by $\Delta_{ji} = \omega_i - \omega_j$. If this detuning is large, both resonant excitation and AC Stark effect are greatly suppressed and the impact of crosstalk on the target qubit becomes negligible (see Section 1.5.1.1). To capture this behaviour, an *effective* crosstalk coefficient β_{ji}^{eff} is introduced, which weights the bare crosstalk coefficient β_{ji} by the spectral overlap between the gate pulse envelope and the detuned transition of the target qubit, $S(\Delta_{ji})$:

$$\beta_{ji}^{\text{eff}} = \beta_{ji} S(\Delta_{ji}) = \beta_{ji} \frac{|\tilde{\Omega}(\Delta_{ji})|}{|\tilde{\Omega}(0)|}, \quad (2.15)$$

where $\tilde{\Omega}(\omega) = \int \Omega(t) e^{-i\omega t} dt$ is the Fourier transform of the gate envelope. In essence, this equation says that the bare crosstalk coefficient β_{ji} only quantifies how much

the drive signal leaks between lines. The target qubit, however, is only sensitive to a portion of the leaked tone near its transition frequency. The effective coefficient β^{eff} directly captures this portion of the crosstalk signal that can contribute to gate errors. In the limit of large detuning ($\Delta \rightarrow \infty$), $\beta^{\text{eff}} \rightarrow 0$, whereas for $\Delta \rightarrow 0$, $\beta^{\text{eff}} \rightarrow \beta$.

For a Gaussian gate pulse of the form $\Omega(t) = \Omega_0 e^{-t^2/2\sigma^2}$, its Fourier transform is given by:

$$\tilde{\Omega}(\omega) = \sqrt{2\pi}\sigma\Omega_0 e^{-\omega^2\sigma^2/2}. \quad (2.16)$$

For a pulse of total duration τ truncated at $\pm 2\sigma$, $\sigma = \tau/4$. Substituting Eq. 2.16 into Eq. 2.15 yields the effective crosstalk coefficient:

$$\beta_{\text{ji}}^{\text{eff}} = \beta_{\text{ji}} e^{-(\Delta_{\text{ji}}\sigma)^2/2}. \quad (2.17)$$

This picture changes when non-Gaussian pulses are used. As shown in Section 1.3.2, transmon qubits are weakly anharmonic and thus susceptible to leakage errors during fast gate operations. The derivative removal by adiabatic gate (DRAG) pulse [66, 72] mitigates this problem by introducing a derivative quadrature component to the conventional Gaussian envelope $\Omega_I(t)$:

$$\Omega(t) = \Omega_I(t) + i\Omega_Q(t) = \Omega_I(t) + i\beta_D \frac{d\Omega_I(t)}{dt}, \quad (2.18)$$

where the DRAG parameter, β_D^1 is the weighting factor for the derivative component. This first-order correction compensates for the effect of the qubit's anharmonicity, thereby suppressing leakage errors and improving gate fidelity. Using the relation $\tilde{\Omega}_I(t) = i\omega\tilde{\Omega}_I(\omega)$, the Fourier transform of a DRAG envelope becomes:

$$\tilde{\Omega}(\omega) = (1 - \beta_D\omega)\tilde{\Omega}_I(\omega). \quad (2.19)$$

From Eqs. 2.16 and 2.19, the effective crosstalk coefficient is therefore:

$$\beta_{\text{ji}}^{\text{eff}} = \beta_{\text{ji}} |1 - \beta_{D,\text{i}}\Delta_{\text{ji}}| e^{-(\Delta_{\text{ji}}\sigma)^2/2}. \quad (2.20)$$

In contrast to the purely Gaussian case, the presence of the derivative term introduces an asymmetry in detuning dependence. As a result, target qubits located above and below the control qubit in frequency see slightly different effective crosstalk even when they have the same detuning magnitude.

2.2 XY crosstalk measurement

There are several experimental approaches for measuring the bare crosstalk coefficient β . One approach is to drive the control qubit at the target qubit's frequency [50], or tune the frequency of the target qubit into resonance with the control qubit [20]. In this case, the leaked tone drives real Rabi oscillations on the target qubit,

¹In the literature, the DRAG parameter is often denoted as α or β . In this thesis, the notation β_D is used to avoid confusion with the qubit's anharmonicity α and the crosstalk coefficient β .

and the observable is the induced Rabi frequency. This forms the basis of the Rabi-like family of crosstalk measurement methods [20, 50, 55, 67, 68] (more on Rabi oscillations can be found in Section A.2.2 of Appendix A).

An alternative approach is to drive the control qubit at its own frequency, resembling the normal gate operating conditions. The leaked field will thus be at the control qubit's frequency, far detuned from the target qubit and does not induce appreciable population transfer. Instead, it mostly produces an AC Stark shift that accumulates as a measurable phase shift on the target qubit. This is the basis of the Ramsey-like [67] and spin-echo [54] families of crosstalk measurement methods (see Sections A.2.3 and A.3.3 of Appendix A for descriptions of the Ramsey and spin-echo sequences, respectively).

Both Rabi- and Stark-shift-based measurements provide the magnitude of the crosstalk coefficient but are insensitive to its phase. To recover the phase, a reference - a second, phase-coherent tone on the target qubit is required, which interferes with the leaked tone. Depending on the implementation, the crosstalk amplitude and phase may be measured in separate experiments [20, 67], or simultaneously in a single one where both the amplitude and phase of the reference tone are swept [54].

2.2.1 Rabi-like measurement family

One example is the method developed by Google Quantum AI to characterise and compensate microwave crosstalk in their 72- and 105-qubit Willow processors [20]. Since the qubits in these processors are frequency-tunable, the target qubit is first tuned to the idle frequency of the control qubit. The control qubit is then driven resonantly while the target qubit is measured (Fig. 2.2a, top). If crosstalk is present, the leaked drive induces Rabi oscillations on the target qubit with frequency Ω_{ctrl} (Fig. 2.2c). To determine the crosstalk amplitude, the target qubit is subsequently driven through its own XY line using the same drive amplitude (Fig. 2.2a, bottom), yielding the direct Rabi frequency Ω_{dir} . The ratio of Ω_{dir} to Ω_{ctrl} gives the magnitude of the crosstalk coefficient, $|\beta|$, which is equivalent to Eq. 2.6.

The crosstalk phase is obtained in a second experiment. The control qubit is again driven resonantly with amplitude A , while a reference tone is simultaneously applied to the target qubit with amplitude $|\beta| \times A$ (Fig. 2.2b). The relative phase between the two tones is swept, producing a Chevron-like interference pattern. The phase at which the induced excitation is nulled corresponds to the maximally destructive interference between the leaked and reference tones (Fig. 2.2d). The crosstalk phase is therefore 180° out of phase with this cancellation phase.

This method is conceptually straightforward and easy to implement. It serves as the basis for the measurement protocol developed in this thesis, although several modifications are required to accommodate our fixed-frequency qubit architecture and existing control system.

2.2.2 Ramsey-like & spin-echo measurement families

The spin-echo approach proposed by Nuerbolati *et al.* [54] is an example of this family and is specifically designed for fixed-frequency qubits (Fig. 2.3a). A spin-echo sequence is applied to the target qubit, while a control pulse is applied to the

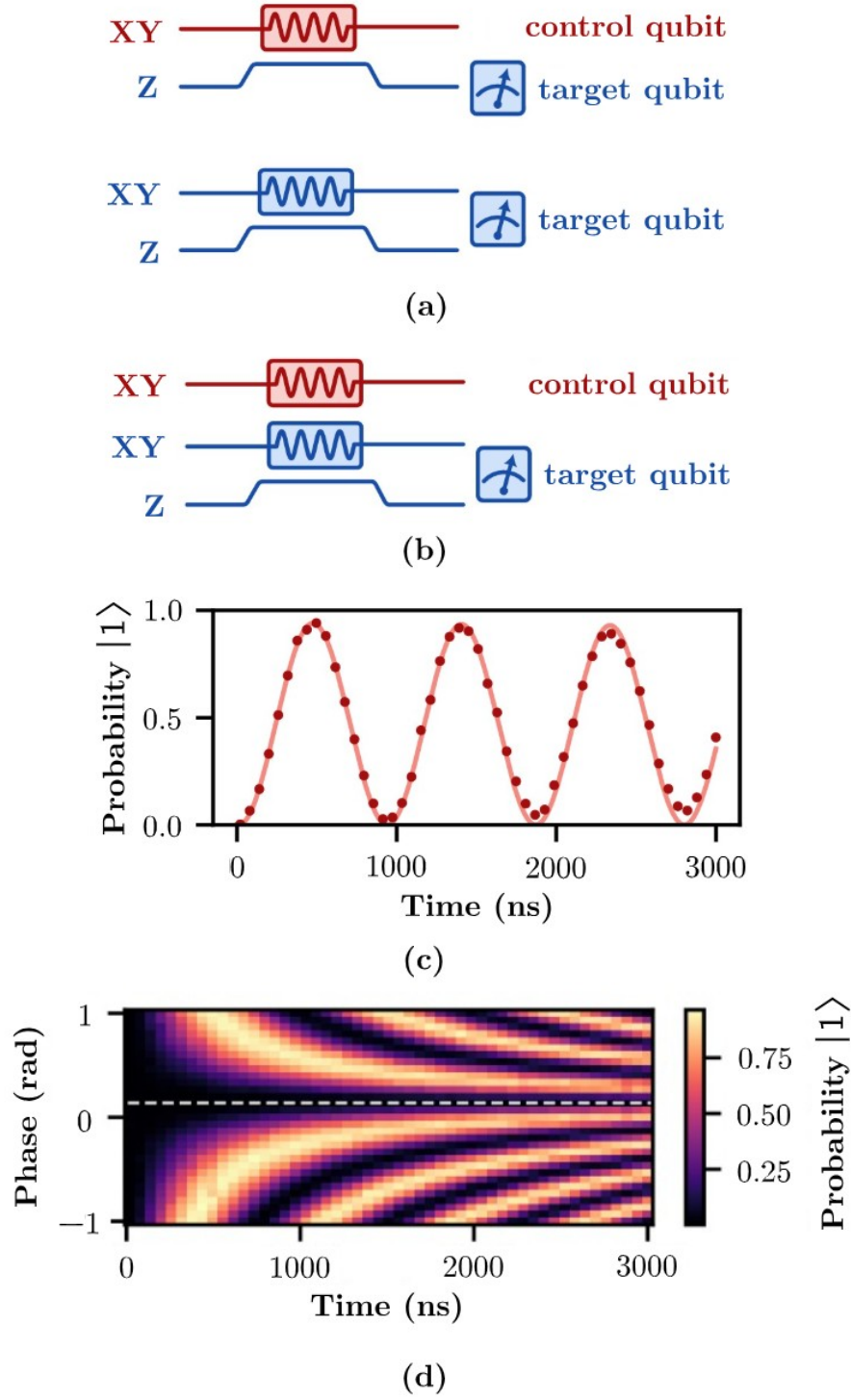


Figure 2.2: A Rabi-like XY crosstalk measurement scheme. (a) Pulse sequence for crosstalk amplitude measurement. (b) Pulse sequence for crosstalk phase measurement. (c) Representative crosstalk-induced Rabi oscillation data corresponding to the first pulse sequence in (a). Ω_{ctrl} is extracted from the sinusoidal fit. (d) Representative phase measurement data corresponding to the pulse sequence in (b). The white dashed line marks the cancellation phase, which is 180° out of phase with the crosstalk phase. Adapted from [20].

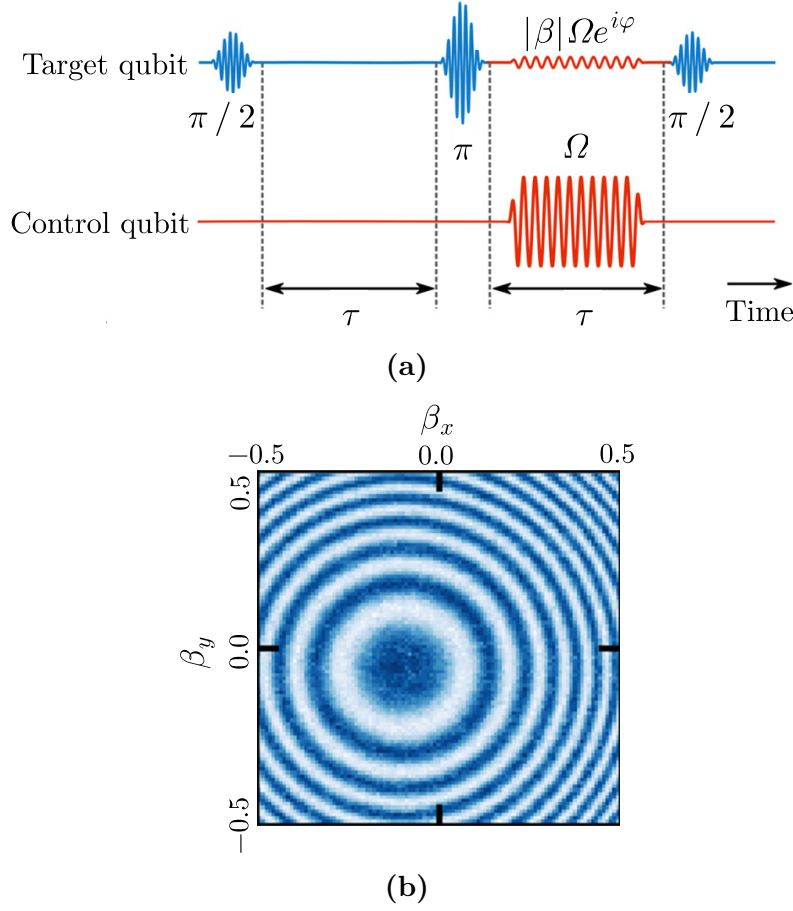


Figure 2.3: A spin-echo XY crosstalk measurement scheme. (a) Pulse sequence for the measurement of both the amplitude and phase of the crosstalk coefficient β . (b) Target qubit's response in the plane of the x- and y-component of the complex crosstalk coefficient β . The centre of the concentric rings marks where the leaked tone is completely cancelled by the compensation tone. Adapted from [54].

control qubit only during the second idle time τ . The phase accumulated from the AC Stark shift is thus not refocused by the echo pulse and can be measured. At the same time, a compensation (reference) tone is applied to the target qubit, with both its amplitude and phase swept simultaneously. Both tones sit at the control qubit's frequency. This produces a concentric interference pattern whose centre marks full cancellation of the leaked signal (Fig. 2.3b). Both the crosstalk amplitude and phase can therefore be obtained from a single two-dimensional measurement. Since the spin-echo sequence suppresses low-frequency noise, the measurement sensitivity can be improved by increasing the free-evolution time τ . The main drawbacks are that the sensitivity is ultimately coherence-limited and the two-dimensional scan is time-consuming.

A related Ramsey-like method was later proposed by Yang *et al.* [67]. Instead of performing a single two-dimensional scan, it measures the crosstalk coefficient through a sequence of one-dimensional measurements, making it more suitable for

large-scale automated calibration. The method uses an $X_{\pi/2}$ Ramsey sequence in the detuned, phase-dominated regime and an X_{π} sequence in the near-resonant regime, providing a unified framework that covers both operating conditions. This approach reduces the measurement time compared to the spin-echo method, but suffers from the lower phase sensitivity at large qubit detunings, requiring additional error-amplification sequences to maintain sufficient visibility.

2.3 XY crosstalk compensation

One of the most commonly used approaches for mitigating XY crosstalk at the signal level is active compensation [54, 67, 68]. The basic idea is to apply an additional microwave counter-tone on the target qubit's XY line that will destructively interfere with the unwanted crosstalk signal leaked from the control qubit before it can drive the target qubit.

Consider again a control qubit i and a target qubit j . The leaked tone from qubit i has the same frequency as the control pulse, f_i , an amplitude $\Omega_{ji} = |\beta_{ji}|\Omega_{ii}$, and a phase shift φ_{ji} relative to the control qubit's intended drive tone. To cancel this leakage, a compensation tone must be applied at the same frequency and amplitude as the leaked tone but with opposite phase, $-\varphi_{ji}$.

This compensation can be expressed naturally using the crosstalk matrix introduced in Section 2.1.1. From Eq. 2.11:

$$\Omega'(t) = C^{-1}\Omega(t), \quad (2.21)$$

where C^{-1} is the inverse of the crosstalk matrix C and is referred to as the compensation matrix. Since the crosstalk coefficients are typically much smaller than 1, C^{-1} can be approximated as:

$$C^{-1} = \frac{1}{1 - \beta_{ji}\beta_{ij}} \begin{pmatrix} 1 & -\beta_{ij} \\ -\beta_{ji} & 1 \end{pmatrix} \approx \begin{pmatrix} 1 & -\beta_{ij} \\ -\beta_{ji} & 1 \end{pmatrix}. \quad (2.22)$$

Let $\Phi'(t)$ and $\Phi(t)$ be two column vectors denoting the total input signals generated by the electronics after compensation has been applied, and the actual signals reaching the qubits, respectively. The two vectors are related by the same crosstalk matrix C :

$$\Phi(t) = C\Phi'(t). \quad (2.23)$$

Since the compensated input signals are obtained by modifying the initial intended pulse scheme,

$$\Phi'(t) = C^{-1}\Omega(t). \quad (2.24)$$

Substituting Eq. 2.24 into Eq. 2.23 gives:

$$\Phi(t) = \Omega(t). \quad (2.25)$$

Thus, after compensation, each qubit receives exactly its intended drive in an ideal scenario where no crosstalk is present.

It is worth noting that the compensation matrix (or crosstalk matrix) is constructed using the bare crosstalk coefficients β , not the effective coefficients β^{eff} .

This is because the crosstalk signal and the compensation pulse both sit at the control qubit's frequency, and thus are equally off-resonant with respect to the target qubit's frequency. Compensation relies only on the electromagnetic coupling between the lines, i.e., how much microwave power leaks from one line to another, rather than by how strongly the target qubit reacts to the off-resonant crosstalk drive.

Extending the discussion to N qubits, the compensation matrix now becomes an $N \times N$ matrix. Many off-diagonal elements of this matrix will be non-zero since XY crosstalk may exist between many qubit pairs. However, in practice, not every crosstalk pathway is worth compensating. As explained in Section 2.1.2, the effect of crosstalk depends on the frequency detuning between the control and target qubit. A qubit with a relatively large β may have a negligible β^{eff} when the leaked tone is sufficiently detuned and contributes insignificantly to the target qubit's gate error. In such cases, actively compensating for the crosstalk signal provides little to no benefit. Furthermore, fully compensating every non-zero off-diagonal elements is undesirable, since it will inject a lot of additional microwave tones which can themselves leak to other qubits and generate further crosstalk. A fully populated compensation matrix also requires measuring a large number of crosstalk coefficients, which is time-consuming and complicated. Additionally, the large number of compensation pulses will place huge demands on the total waveform memory and require a sufficient number of sequencers available for each qubit control module.

One practical solution is therefore to first prune the compensation matrix based on frequency detuning. The relevant frequency scale is set by the inverse of the gate duration, $1/\tau$, which characterises the spectral bandwidth of the gate pulse. For qubit pairs with frequency detuning much larger than $1/\tau$, their effective crosstalk is expected to be negligible regardless of the magnitude of their electromagnetic coupling. Their corresponding matrix elements can therefore be safely set to zero and no compensation is needed. In contrast, pairs for which the detuning is comparable to or smaller than $1/\tau$ should be investigated further.

This frequency-based criterion leaves only a small subset of potentially relevant qubit pairs, which can then be evaluated based on their effective crosstalk coefficient β^{eff} . If a pair satisfies

$$|\beta^{\text{eff}}| < \beta^{\text{th}}, \quad (2.26)$$

where β^{th} is a certain threshold that sets the single-qubit gate error budget, the corresponding matrix element can again be safely pruned to zero. Only qubit pairs with $|\beta^{\text{eff}}| > \beta^{\text{th}}$ are retained and compensated. In this way, compensation is limited to crosstalk pathways that are expected to exhibit a measurable impact on gate performance, while minimising the number of compensation pulses applied and the associated experimental overhead.

An alternative criterion for pruning the compensation matrix is to directly assess the impact of crosstalk on gate performance by comparing the single-qubit gate error obtained from simultaneous operation of a qubit pair with that measured during isolated operation. Only pairs that show a significant increase in gate error are then characterised for crosstalk and compensated. Further studies are needed to verify if there is a correlation between β^{eff} and the increase in gate error, and to

determine which of the two criteria is more reliable and practical for pruning the compensation matrix. For larger processors, additional criteria may be required as the number of frequency collisions and crosstalk pathways increases.

To conclude this chapter, it is important to again distinguish the roles of the two crosstalk coefficients. The bare crosstalk coefficient β defines the amplitude and phase of the compensation pulse required for crosstalk cancellation, whereas the effective crosstalk coefficient β^{eff} quantifies the crosstalk level experienced by the target qubit after accounting for frequency detuning and therefore may be used to decide if compensation is needed in the first place.

3

Experimental Setup

This chapter provides an overview of the quantum device and measurement setup used throughout this thesis. Section 3.1 introduces the architecture of the 25-qubit flip-chip quantum processor developed in the QC2 group. The qubit frequency allocation scheme adopted in the chip, as well as fabrication-induced frequency variations, is discussed in detail in Section 3.2, as these topics are closely related to frequency collisions and microwave crosstalk. The cryogenic measurement setup is described in Section 3.3, while Section 3.4 presents the software framework developed within the QC2 group for the automated calibration, characterisation, and benchmarking of superconducting qubits. Further details on the chip design and experimental setup can be found in previous publications from the group [47, 50].

3.1 The 25-qubit flip-chip quantum processor

The architecture of the 25-qubit quantum processor used in this work is shown in Fig. 3.1. The device follows the flip-chip architecture described in Section 1.4.2, where two chips - a Q-chip and a C-chip are bonded together via superconducting indium bumps. The 25 fixed-frequency transmon qubits are arranged in a 5×5 square lattice, with five qubits per row and five rows in total. The qubits are labelled from q01 to q25, from left to right and top to bottom, as illustrated in Fig. 3.2. The nearest-neighbour spacing between the qubits is 2 mm.

Each qubit is capacitively coupled to a quarter-wavelength coplanar waveguide resonator, which is in turn coupled to a feedline for dispersive readout. The five resonators within each row share a common feedline, resulting in a total of five feedlines across the processor. Qubit control is possible through each qubit's dedicated XY line. To enable two-qubit gates, the qubits employ an Xmon geometry, allowing them to be coupled to all nearest neighbours via tunable couplers. Depending on the qubit's location within the lattice, it can be coupled to two, three, or four other qubits. The couplers are flux-tunable via their dedicated Z lines.

3.2 Qubit frequency allocation scheme

As explained in Section 1.5.1.1, the effect of XY crosstalk is strongly suppressed when nearby qubits are far detuned from one another. This highlights the importance of a qubit frequency allocation scheme in the chip design that maximises frequency separations between neighbouring qubits whenever possible.

One such scheme is shown in Fig. 3.2, where the first transition frequencies f_{01}

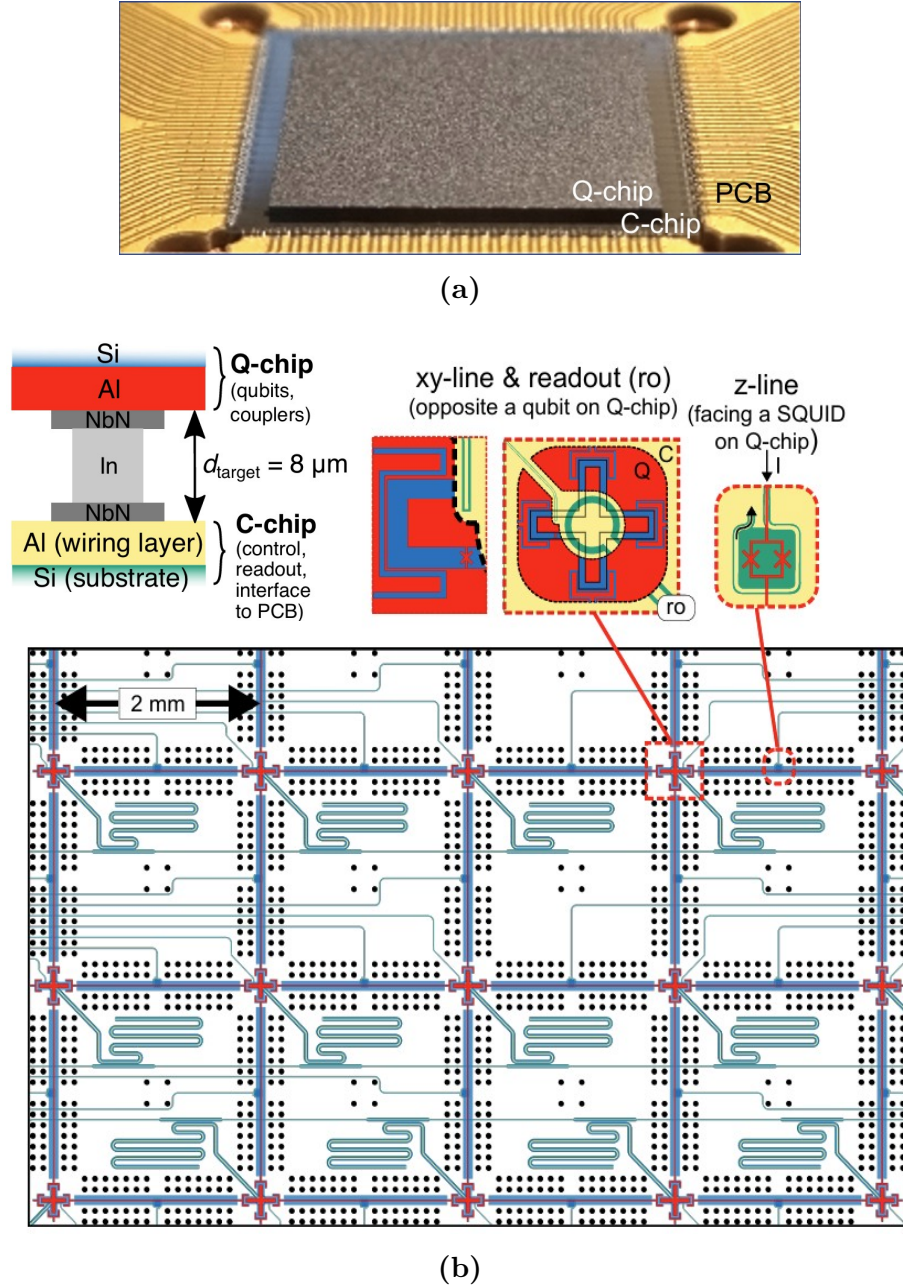


Figure 3.1: The 25-qubit flip-chip quantum processor. (a) An image of the quantum processor. (b) Top left: Cross-sectional schematic of the material stack in the flip-chip quantum processor. The Q-chip hosts the qubits and couplers, while the C-chip contains the control and readout circuitry and the interface to the PCB. Top middle: An Xmon qubit on the Q-chip capacitively coupled to a readout resonator and an XY line on the C-chip. Top right: The SQUID loop of a tunable coupler and its corresponding Z line. Bottom: A layout of part of the chip, showing qubits q11 - q25 (red crosses) in the last three rows, their dedicated couplers, transmission lines, readout resonators, control lines, and indium bumps (black dots). Adapted from [50].

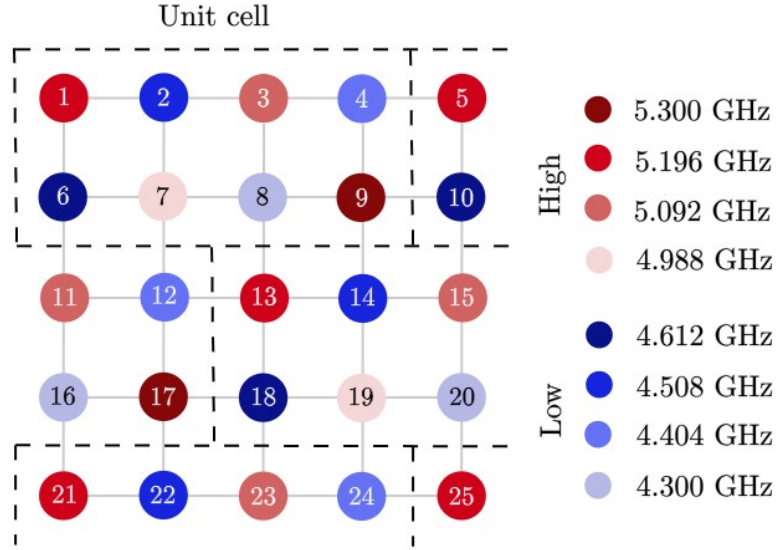


Figure 3.2: Qubit frequency allocation scheme in the 25-qubit device. The first transition frequencies f_{01} of the qubits are arranged in a 5×5 square grid in an alternating fashion between two frequency groups: high (red) and low (blue). Each group consists of four different frequency values, illustrated by different hues of red/blue. The grey solid lines represent the couplers connecting the qubits, while the black dashed lines represent the boundaries of the frequency unit cells, each of which contains all eight possible frequency values from both groups.

of the qubits are allocated in an alternating pattern of *high* and *low* frequencies [47, 73]. The high-frequency group (red) consists of four different target frequencies between 4.988 and 5.300 GHz, all sharing a common anharmonicity of -260 MHz. The low-frequency group (blue) comprises four other frequencies ranging from 4.300 to 4.612 GHz, with an anharmonicity of -156 MHz. This arrangement makes sure that nearest-neighbour qubits belong to different frequency groups, and two-qubit gates can only be implemented between a high-frequency and a low-frequency qubit. Together, these eight values of qubit frequencies form a unit cell that is repeated across the processor lattice.

Since only eight distinct frequencies are used across the 25-qubit device, some qubits inevitably share the same target frequency. To minimise the risk of XY crosstalk between these qubits, they must be separated as far as possible within the chip layout. This is achieved by shifting the frequency unit cells by half of their length. As a result, the shortest centre-to-centre distance between two qubits with the same target frequency is approximately 5.7 mm. This corresponds to a distance of 8 mm when measuring along the shortest path through the coupler network. Furthermore, to reduce the probability of unwanted electromagnetic coupling within the microwave circuitry and hence XY crosstalk via this channel (Fig. 1.8), the spacings between XY lines, as well as between themselves and the qubits, are also optimised. For example, the XY lines corresponding to qubits with the same frequency are routed to different regions of the chip.

This frequency allocation strategy is simple, scalable, and tileable, making it

suitable for larger processor layouts. However, it is not without limitations. In particular, the scheme primarily considers $f_{01} - f_{01}$ collisions and does not explicitly account for $f_{01} - f_{12}$ collisions. As a result, most of the XY crosstalk effects observed in the 25-qubit device originate from the latter.

However, even with a delicately designed frequency allocation scheme, the actual qubit frequencies almost always deviate from their target values due to fabrication-induced parameter variations. The reproducibility of the qubit frequency is mainly limited by the JJ, whose small dimensions (compared to the much larger shunted capacitor) make it particularly sensitive to lithographic variations [47]. The two dominant sources of error are fluctuations in the junction area and non-uniformities in the thickness of the oxide tunnel barrier [47].

These frequency variations can have both beneficial and detrimental consequences. On one hand, they may reduce the intended detuning between two qubits, making them more susceptible to crosstalk. On the other hand, they can also lift the frequency degeneracy between two qubits with the same design frequency, leading to a suppression of crosstalk. For example, by design, qubits q06 and q10 share the same $|0\rangle - |1\rangle$ transition frequency of 4.612 GHz. However, due to fabrication-induced variations, their measured frequencies are approximately 4.621 GHz and 4.434 GHz, respectively. This large frequency detuning of 187 MHz sufficiently suppresses crosstalk, resulting in no observable impact on their individual single-qubit gate fidelities.

One solution to improve frequency reproducibility is to fabricate larger junctions with thicker oxide barriers to lower their sensitivity to fabrication variations [73]. Another approach is post-fabrication electrical tuning of the junction, in which voltages are applied to controllably modify the normal-state resistance of the junction, thereby changing its Josephson energy and allowing the qubit frequency to be tuned towards its target value after fabrication [74]. The QC2 group is investigating both of these approaches.

3.3 Cryogenic measurement setup

Superconducting qubits must be operated in the tens-of-millikelvin temperature range to ensure that the circuits remain superconducting and thermal excitations are strongly suppressed. This is achieved using a Bluefors dilution refrigerator (or cryostat), whose outermost vacuum chamber thermally isolates the interior from the exterior environment. The cryostat consists of a series of progressively colder stages, from room temperature (300 K) down to approximately 10 mK. Each stage is enclosed by a thermally anchored shield that blocks the blackbody radiation emitted from the warmer stages above. The 25-qubit quantum processor, after being packaged, is mounted on the mixing-chamber stage of the refrigerator at the base temperature of 10 mK (Fig. 3.3).

The chip is connected to the room-temperature Qblox clusters (Fig. 3.4) via coaxial control and readout lines. These clusters generate microwave pulses used to operate the device and also receive the corresponding readout signals. Due to limitations on the number of modules available per cluster, two separate clusters, labelled A and B, are used to control the qubits (Fig. 3.4a). Cluster A controls

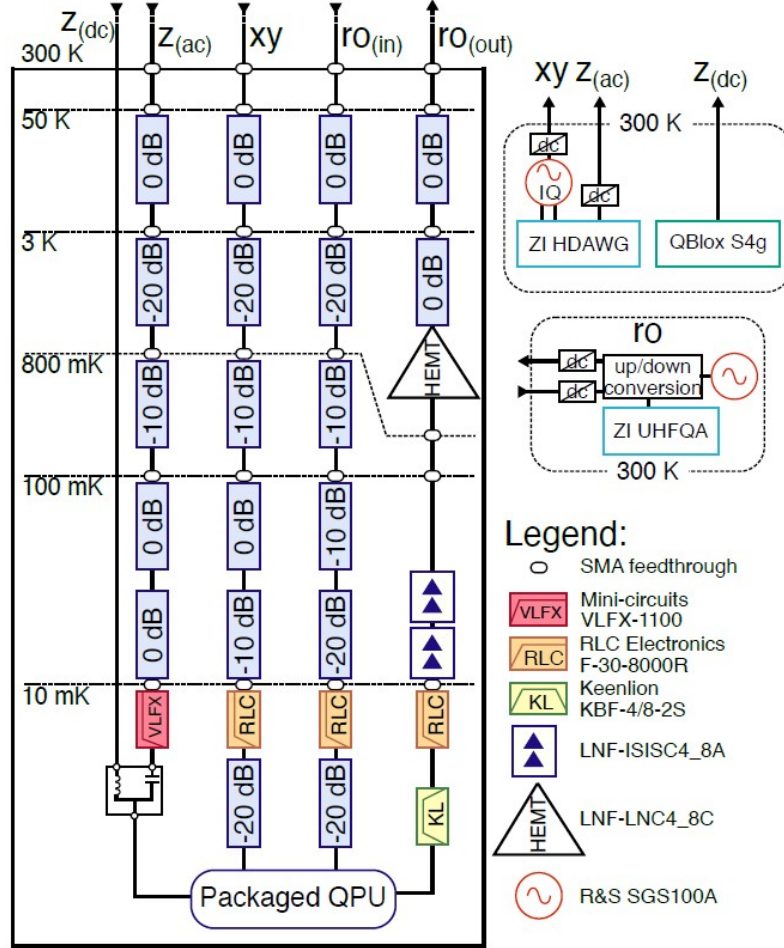


Figure 3.3: Wiring diagram for the measurement setup at room and cryogenic temperatures. The 25-qubit flip-chip processor is mounted on the mixing-chamber stage (10 mK) of a Bluefors dilution refrigerator. Control and read-out signals are routed from room-temperature Qblox control electronics through coaxial lines that are thermally anchored and attenuated (light blue rectangles) at successive stages to suppress thermal noise and infrared radiation. The weak signal coming off the readout-output line is amplified by a high electron mobility transistor (HEMT) and protected from back-propagating noise by a chain of isolators (dark blue triangles). Low-pass filters (VLFX and RLC) and a band-pass filter (KL) reject out-of-band noise. Figure from [50].

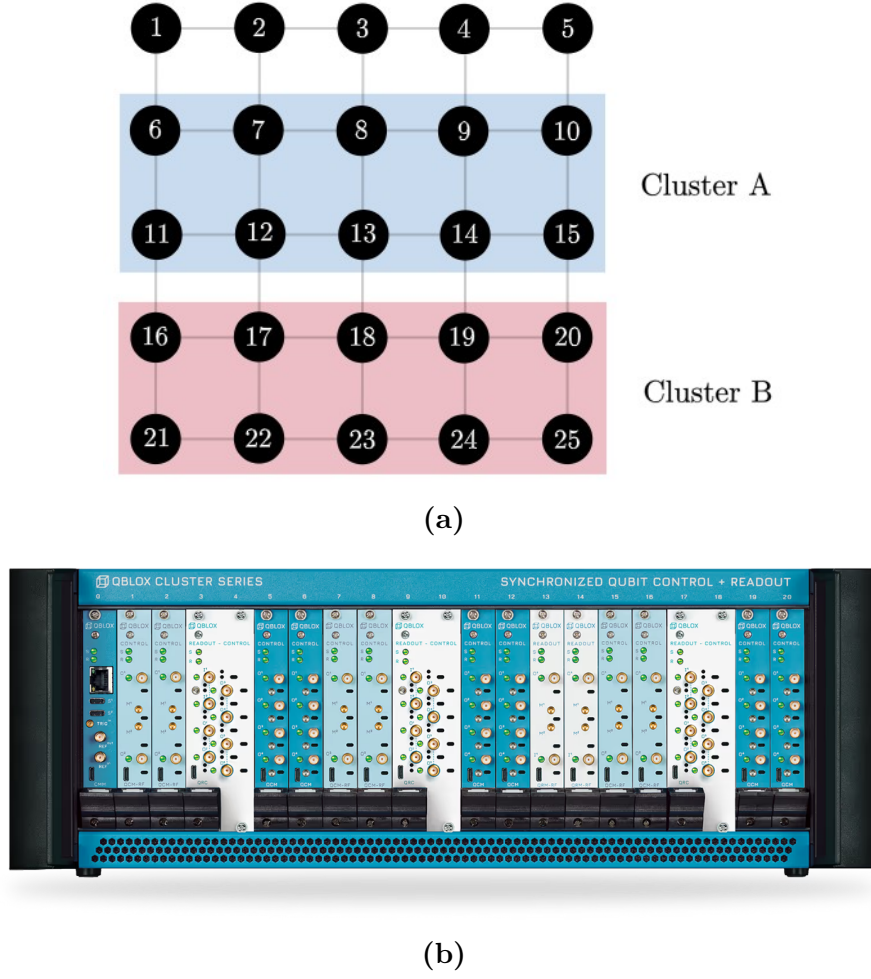


Figure 3.4: Qblox clusters. (a) The qubits are controlled by two separate clusters: qubits q06 - q15 by cluster A, and qubits q16 - q25 by cluster B. Qubits q01 - q05 are not connected. (b) An example of a Qblox cluster with different types of modules (blue and white columns), each with a certain number of input and output channels. QCM-RF and QRM-RF modules are used in our setup for qubit control and readout respectively. Image from [75].

qubits q06 - q15, while cluster B controls qubits q16 - q25. The five qubits in the first row (q01 - q05) are not currently connected due to wiring constraints. As a result, although the chip has 25 qubits, only 20 are accessible with the present measurement setup. For operational reasons, experiments are limited to a single cluster, so at most 10 qubits at a time. Operating both clusters together is however possible and has already been demonstrated in previous work by the group with no known limitations [47]. Therefore, extending the methods developed in this thesis to the full chip is straightforward.

Pulses applied to the XY and $Z_{(AC)}$ lines are generated by QCM-RF II modules, while the readout tones are generated and measured using QRM-RF modules. Each module carries its own local oscillator (LO) and can produce microwave signals between 2 and 18.5 GHz. DC bias currents on the $Z_{(DC)}$ lines are supplied by Qblox



Figure 3.5: The experimental setup. A photo of the cryogenic measurement setup showing the Bluefors dilution refrigerator (right), Qblox clusters (left), other electronics, and cables.

SPI racks (S4g).

Since the control and readout-input lines connect directly to room-temperature electronics, thermal noise can propagate to the quantum processor and induce qubit decoherence. To suppress this noise, the incoming lines are attenuated and filtered at several stages throughout the refrigerator. The attenuators dissipate the incoming signals together with their accompanying thermal photons as heat. Therefore, they must be distributed across different stages to avoid excessive heating at the mixing chamber. Low-pass filters (Mini-Circuits VLFX-1100 and RLC F-30-8000R) installed at the mixing-chamber stage further suppress high-frequency noise and infrared radiation on the incoming lines.

The readout signals emitted by the qubits are extremely weak and thus need to be amplified with minimal added noise before they can be detected at room temperature. This is done by a low-noise HEMT amplifier at 3 K. Between the device and the amplifier, in addition to an RLC low-pass filter, a Keenlion KBF-4/8-2S band-pass filter is inserted to transmit only the 4 - 8 GHz readout band. Moreover, a chain of cryogenic isolators (LNF-ISISC4_8A) is used as a one-way valve, allowing upward-travelling readout signals to pass freely while preventing noise and reflections from propagating back down and disturbing the qubits. After amplification, the signals return to the QRM-RF modules, where they are down-converted to an intermediate frequency (IF) in the hundreds-of-megahertz range and digitally integrated to recover the qubit state.

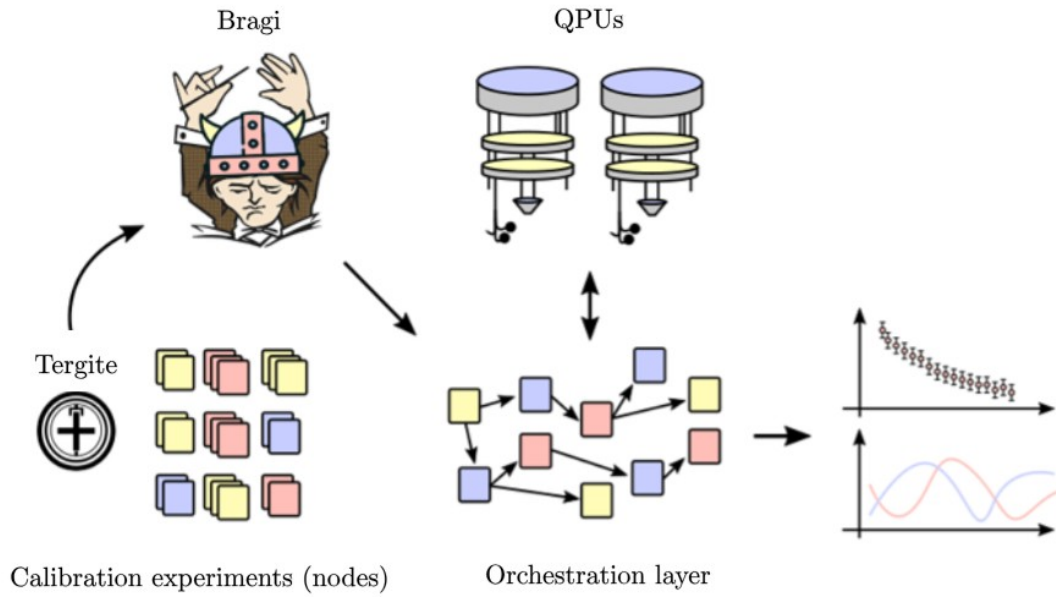


Figure 3.6: The software framework. Tergite Autocalibration consists of a calibration supervisor and a library of calibration experiments implemented as nodes. Bragi orchestrates different layers of a qubit calibration routine or other complex sequences of experiments that quantify the performance of the quantum processor. Graphics adapted from M. Robbiati.

3.4 The software stack

The calibration and characterisation of the 25-qubit quantum processor is performed using an open-source Python-based software framework developed by the QC2 group at Chalmers University of Technology (Fig. 3.6). It consists of two packages: Tergite Autocalibration and Bragi.

Tergite Autocalibration [76] is an operational library of calibration and characterisation experiments. Each calibration experiment is designed to measure a specific quantity of interest (QOI), i.e., a parameter of the quantum processor, such as qubit and resonator frequencies, π pulse amplitudes, and CZ phase. Each of them is implemented as a *node* consisting of a measurement schedule and a post-processing analysis scheme, which together extract the relevant QOI. In addition, the package includes characterisation experiments, such as T_1 , T_2 , and RB, which are used to evaluate qubit coherence and gate performance.

The other repository, Bragi, acts as an orchestrator that controls the overall experimental pipeline, such as the concatenation of experiments needed to calibrate single-qubit gates described in Section A.5 of Appendix A. Bragi also handles complex workflows that are difficult to implement within Tergite Autocalibration itself, including time correlation and crosstalk studies.

Communication between the software stack and the Qblox hardware (clusters and SPI racks) is handled through the Python-based Quantify Scheduler framework. A Redis database is used to store the calibration status and QOIs and is updated after each experiment.

4

XY Crosstalk: Methods, Results & Discussion

The XY crosstalk measurement scheme developed in this work is a Rabi-like approach based on Google Quantum AI’s protocol [20], which is described in Section 2.2.1. It is modified to better suit the fixed-frequency qubits in the 25-qubit processor and the calibration experiments already established in the software stack. The crosstalk experiments are implemented as four separate nodes: XY Crosstalk Amplitude 01, XY Crosstalk Amplitude 12, XY Crosstalk Phase 01, and XY Crosstalk Phase 12. As their names suggest, the amplitude nodes determine the amplitude of the crosstalk coefficient $|\beta|$, while the phase nodes measure its phase φ . $f_{01} - f_{01}$ and $f_{01} - f_{12}$ crosstalk are characterised separately.

4.1 XY crosstalk amplitude measurement

The pulse sequence used for measuring the amplitude of the crosstalk coefficient is shown in Fig. 4.1a.

For $f_{01} - f_{01}$ crosstalk, both qubits are first reset and prepared in their ground state. A Gaussian pulse is then applied to the control qubit i , but at the target qubit j ’s first transition frequency, f_j^{01} , before the target qubit’s response is measured. The pulse duration is kept fixed while the pulse amplitude is swept. Since the leaked tone arriving at the target qubit is resonant with its $|0\rangle - |1\rangle$ transition, it drives coherent Rabi oscillations on the target qubit. From the oscillations, the induced π pulse amplitude, A_{ji}^π , can be extracted (Fig. 4.2). Physically, this can be interpreted as the pulse amplitude required to be applied on the control qubit to induce a π rotation on the target qubit through XY crosstalk.

The bare crosstalk coefficient amplitude is obtained by comparing this induced π pulse amplitude with the calibrated π pulse amplitude of the control qubit, A_{ii}^π , which has been measured previously from a standard Rabi oscillation experiment (see Section A.2.2):

$$|\beta_{ji}| = \frac{A_{ii}^\pi}{A_{ji}^\pi}. \quad (4.1)$$

This expression is equivalent to the definition given by Eq. 2.6 in Section 2.1.1. To prove this, assume a linear response between the applied pulse amplitude A and the resulting Rabi rate Ω for both the direct drive on the control qubit i and the crosstalk drive on qubit j :

$$\Omega_{ii} = \kappa_{ii} A, \quad (4.2)$$

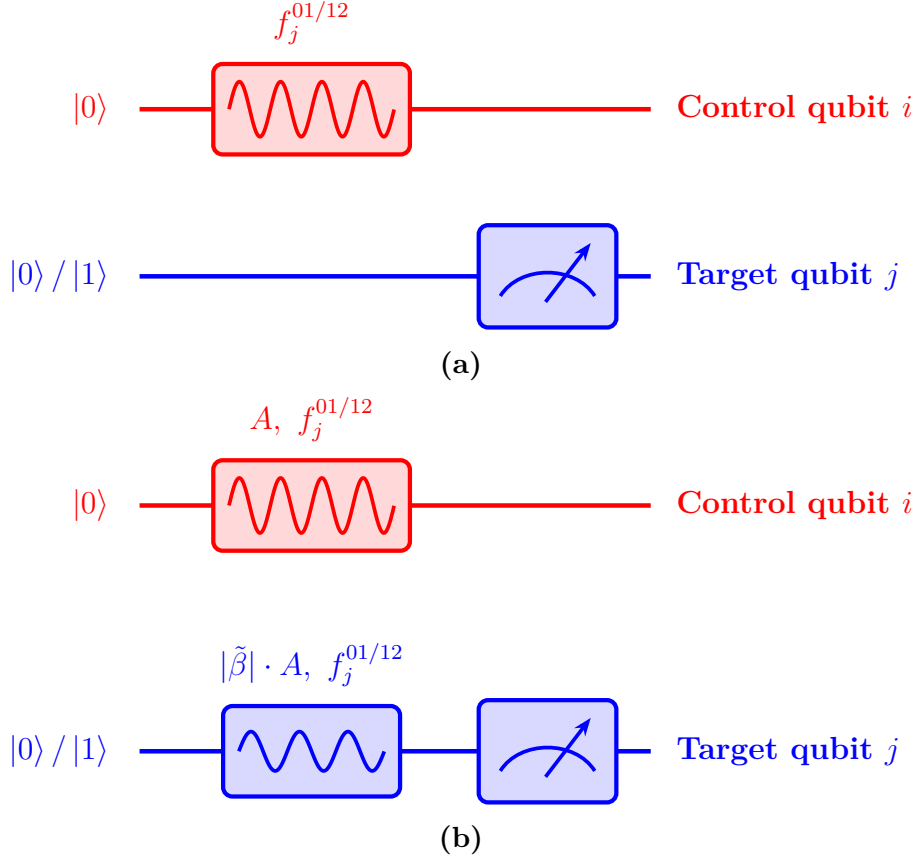


Figure 4.1: Pulse sequences for XY crosstalk measurement. (a) Amplitude measurement: The control qubit is driven by a Gaussian pulse at the target qubit's frequency with fixed duration and varying amplitude. The target qubit's response is then measured. (b) Phase measurement: Both qubits are driven at the target qubit's frequency simultaneously, each by a Gaussian pulse with varying duration and phase. The amplitude of the target pulse is scaled by a factor of $|\tilde{\beta}|$ compared to the control pulse amplitude to enable complete cancellation of the two tones. Finally, the target qubit's response is measured.

$$\Omega_{ji} = \kappa_{ji}A, \quad (4.3)$$

where κ_{ii} and κ_{ji} describe the linear conversion from the applied amplitudes to the corresponding Rabi rates. For any fixed amplitude A , the bare crosstalk coefficient is therefore:

$$|\beta_{ji}| = \frac{\Omega_{ji}}{\Omega_{ii}} = \frac{\kappa_{ji}}{\kappa_{ii}}. \quad (4.4)$$

For a Gaussian pulse with envelope $s(t)$, duration τ , and amplitude A , the accumulated rotation angle is determined by the pulse area, or equivalently by the time integral of the instantaneous Rabi rate:

$$\theta = \kappa \cdot A \cdot \int_0^\tau s(t)dt. \quad (4.5)$$

For a π pulse, the rotation angle must satisfy $\theta = \pi$. Eq. 4.5 then becomes:

$$\pi = \kappa_{ji} \cdot A_{ji}^\pi \cdot \int_0^{\tau_{ji}} s(t)dt, \quad (4.6)$$

$$\pi = \kappa_{ii} \cdot A_{ii}^\pi \cdot \int_0^{\tau_{ii}} s(t)dt. \quad (4.7)$$

If the same pulse shape and duration are used in both measurements, the integrals in Eqs. 4.6 and 4.7 are identical, and therefore:

$$\frac{A_{ii}^\pi}{A_{ji}^\pi} = \frac{\kappa_{ji}}{\kappa_{ii}} = \frac{\Omega_{ji}}{\Omega_{ii}} = |\beta_{ji}|, \quad (4.8)$$

uniting Eqs. 4.1, 4.4, and 2.6. It is important to note that this equation is only valid when the same pulse duration is used in both the standard Rabi oscillation and the crosstalk amplitude measurements. In practice, however, since the leaked tone is much weaker than the direct drive, a pulse with a much longer duration (τ_{ji}) than the standard gate duration (τ_{ii}) is usually used in the crosstalk amplitude measurement. This longer duration increases the number of Rabi rotations observable within the limited window of pulse amplitudes swept, allowing A_{ji}^π to be extracted with higher precision. In this case, the two π pulse amplitudes must therefore be normalised by their corresponding pulse durations, giving:

$$|\beta_{ji}| = \frac{A_{ii}^\pi \cdot \tau_{ii}}{A_{ji}^\pi \cdot \tau_{ji}}. \quad (4.9)$$

This experiment relies on several parameters that are already obtained from the single-qubit calibration chain: readout frequency and amplitude to measure the target qubit's response, control qubit's first transition frequency to perform Rabi oscillations, and control qubit's π pulse amplitude for calculating $|\beta|$.

The $f_{01} - f_{12}$ crosstalk measurement follows the same principle. However, the target qubit is now first excited to state $|1\rangle$ by applying an X gate, and the control tone is applied at the target qubit's second transition frequency, f_j^{12} , instead of f_j^{01} . The resulting Rabi oscillations are analysed in the same way as for $f_{01} - f_{01}$ crosstalk.

Compared with Google's implementation, the method in this thesis differs in two respects. First, in Google's protocol, the target qubit is tuned to the idle frequency

of the control qubit. This is possible because their qubits are flux-tunable. Since the qubits used in this thesis are fixed-frequency qubits, the same resonance condition is instead achieved by driving the control qubit directly at the target qubit’s frequency. Second, Google extracts the crosstalk coefficient using time-swept Rabi oscillations, whereas in this work, the amplitude is swept instead. This approach is adopted because the existing Rabi oscillation experiment in **Tergite Autocalibration** is performed in this way. The two approaches are physically equivalent.

An example of $f_{01} - f_{12}$ crosstalk amplitude results obtained for qubits q07 and q09 is shown in Fig. 4.2. The f_{01} transition of q07 is separated from the f_{12} transition of q09 by only approximately 3.6 MHz, making this pair particularly susceptible to $f_{01} - f_{12}$ crosstalk. The extracted bare crosstalk coefficients are comparable in the two directions, with $|\beta| = 0.0273$ (−31 dB) from q07 to q09 and $|\beta| = 0.0169$ (−35 dB) from q09 to q07.

Despite this symmetry in the electromagnetic coupling between the two lines, only q09 experiences significant crosstalk. Since its f_{12} transition lies very close to the f_{01} transition of q07, the effective crosstalk coefficient is approximately 96.7 % of the bare value. This results in a 1.6-fold increase in the average single-qubit gate error, from 0.073 % when RB is performed in isolation to 0.12 % when measured simultaneously with q07 (Figs. 4.4 and 4.6a). Additionally, qubit q09’s leakage rate increases by almost ten times when it is driven simultaneously with qubit q07 (Fig. 4.4 and 4.7a). More details on how the average single-qubit gate error, single-qubit gate fidelity, and leakage rate are extracted from RB are given in Section A.3.4.

The opposite direction shows a very different behaviour. The leaked f_{01} drive from q09 is detuned from the relevant transitions of q07 by several hundred megahertz, suppressing the effective crosstalk coefficient to essentially zero. Consequently, neither the gate fidelity nor the leakage rate of q07 changes measurably when RB is performed simultaneously with q09 (Figs. 4.6a, 4.7a, and B.3), and thus no active compensation is required. This directional asymmetry is characteristic of $f_{01} - f_{12}$ crosstalk.

In contrast, $f_{01} - f_{01}$ crosstalk is typically more symmetric, with both qubits being similarly affected by microwave leakage from the other. An example of $f_{01} - f_{01}$ crosstalk amplitude results obtained for qubits q11 and q15 is shown in Fig. B.1 in Appendix B.

4.2 XY crosstalk phase measurement

The pulse sequence used to measure the phase of the crosstalk coefficient is shown in Fig. 4.1b.

For $f_{01} - f_{01}$ crosstalk, both qubits are first initialised in state $|0\rangle$. Two Gaussian pulses are then applied simultaneously at the target qubit’s first transition frequency, f_j^{01} - one to the control qubit with fixed amplitude A , and the other to the target qubit with amplitude $|\tilde{\beta}| \cdot A$. The pulse duration and the relative phase between the two signals are swept. Finally, the target qubit’s response is measured, forming a Chevron-like interference pattern as a function of pulse duration and relative phase. The cancellation phase, $\varphi_{ji}^{\text{canc}}$, is the phase at which the target tone cancels the

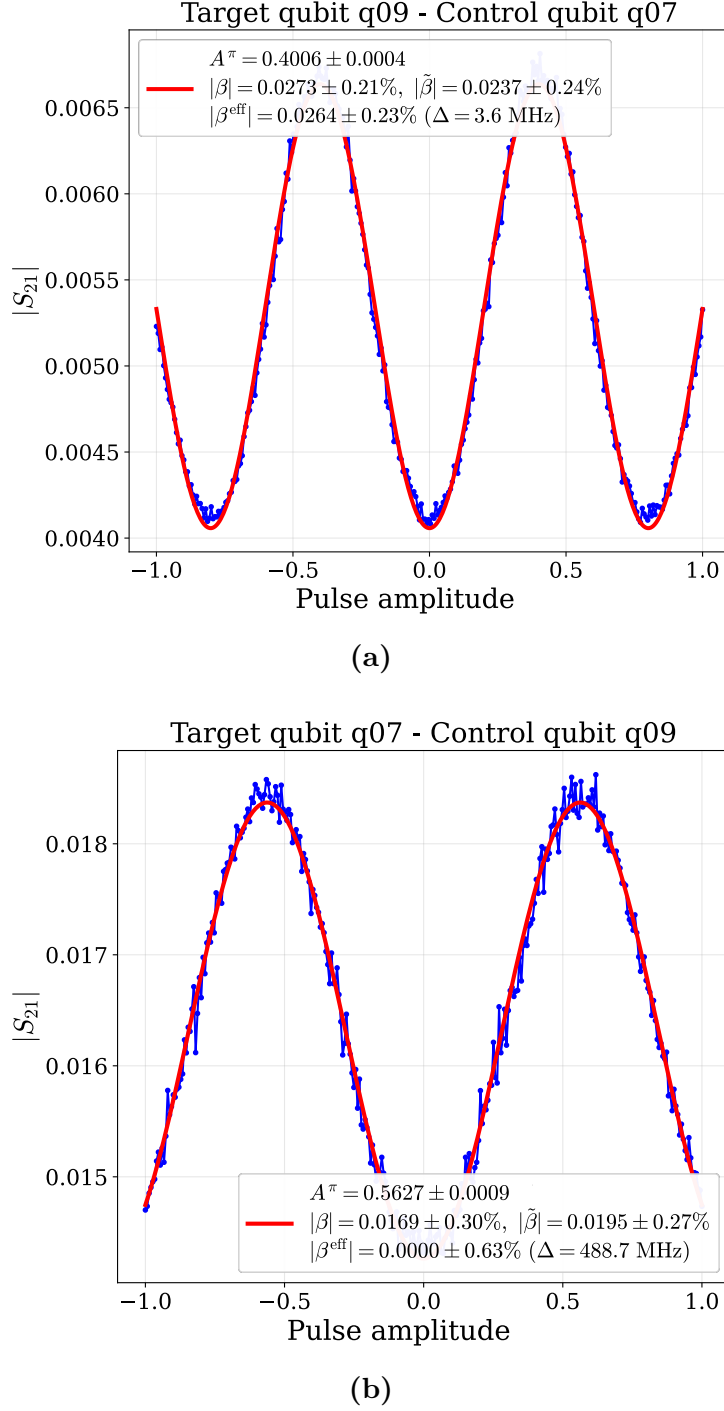


Figure 4.2: Representative $f_{01} - f_{12}$ crosstalk amplitude results for qubits q07 and q09. (a) Qubit q07 acts as the control qubit and q09 as the target qubit. (b) Qubit q09 acts as the control qubit and q07 as the target qubit. Although the bare crosstalk coefficients $|\beta|$ are of similar amplitude in both directions, the effective crosstalk is strongly suppressed in (b) due to the large detuning between the leaked drive at f_{01} of q09 and the f_{12} transition of q07. $|\beta|$ is the amplitude scaling factor used in the crosstalk phase measurement (see Section 4.2).

crosstalk tone perfectly for all pulse durations. The crosstalk phase, φ_{ji} , is thus:

$$\varphi_{ji} = \varphi_{ji}^{\text{canc}} - 180^\circ. \quad (4.10)$$

The parameter $|\tilde{\beta}|$ is an amplitude scaling factor used to set the strength of the reference tone applied to the target qubit. It is chosen such that the reference tone has the same drive strength as the crosstalk tone, allowing the two signals to cancel completely at the cancellation phase. In such condition, a well-defined Chevron pattern will be produced with high contrast between bright and dark fringes. If the two strengths are mismatched, a residual net drive remains even at the cancellation phase. For sufficiently long pulses, this residual drive accumulates into multiple π rotations, causing the nominally dark centre of the Chevron pattern to oscillate. For a pair of control qubit i and target qubit j , the scaling factor $|\tilde{\beta}_{ji}|$ is given by:

$$|\tilde{\beta}_{ji}| = \frac{A_{jj}^\pi \cdot \tau_{jj}}{A_{ji}^\pi \cdot \tau_{ji}}. \quad (4.11)$$

It is important to distinguish $|\tilde{\beta}_{ji}|$ from the amplitude of the bare crosstalk coefficient $|\beta_{ji}|$ defined in Eq. 4.9. $|\beta_{ji}|$ quantifies the strength of the leaked signal relative to the intended drive on the control qubit and is therefore referenced to the control qubit's calibrated π pulse amplitude, A_{ii}^π . In contrast, $|\tilde{\beta}_{ji}|$ is used solely in the phase measurement to generate a reference pulse on the target qubit with a drive strength matched to that of the leaked tone. It is therefore referenced to the target qubit's calibrated π pulse amplitude, A_{jj}^π , rather than that of the control qubit.

The $f_{01} - f_{12}$ crosstalk measurement follows the same principle. However, the target qubit is now first prepared in state $|1\rangle$ by applying an X gate, and both pulses are applied at the target qubit's second transition frequency, f_j^{12} . The resulting Chevron pattern is analysed identically to extract the crosstalk phase.

An example of $f_{01} - f_{12}$ crosstalk phase results obtained for qubits q07 and q09 is shown in Fig. 4.3. Rather than sweeping both the relative phase and pulse duration in a full two-dimensional map like as in Google's implementation, this experiment is more sufficiently implemented as three separate one-dimensional phase sweeps at selected pulse durations. These few slices capture sufficient information to extract the cancellation phase with high angular resolution, while significantly reducing both the acquisition time and the total waveform memory required.

At first glance, the interference pattern in Fig. 4.3a appears to contain a second cancellation phase around -100° . This feature, however, is not an additional cancellation point, but the boundary between two consecutive 360° cycles of the Chevron pattern. The measurement is always performed over a fixed phase range from -180° to 180° . This range does not necessarily coincide with one complete Chevron pattern, which repeats every 360° . Depending on the actual cancellation phase, the scan can therefore contain the end of one cycle followed by the beginning of the next, with the two parts meeting at a boundary within the measurement range. The same boundary is visible near 25° in Fig. 4.3b. This boundary is important because it corresponds to the phase of maximum constructive interference between the leaked tone and the reference tone. Therefore, it directly gives the crosstalk phase, φ_{ji} , which is 180° out of phase with the cancellation phase, $\varphi_{ji}^{\text{canc}}$. At this

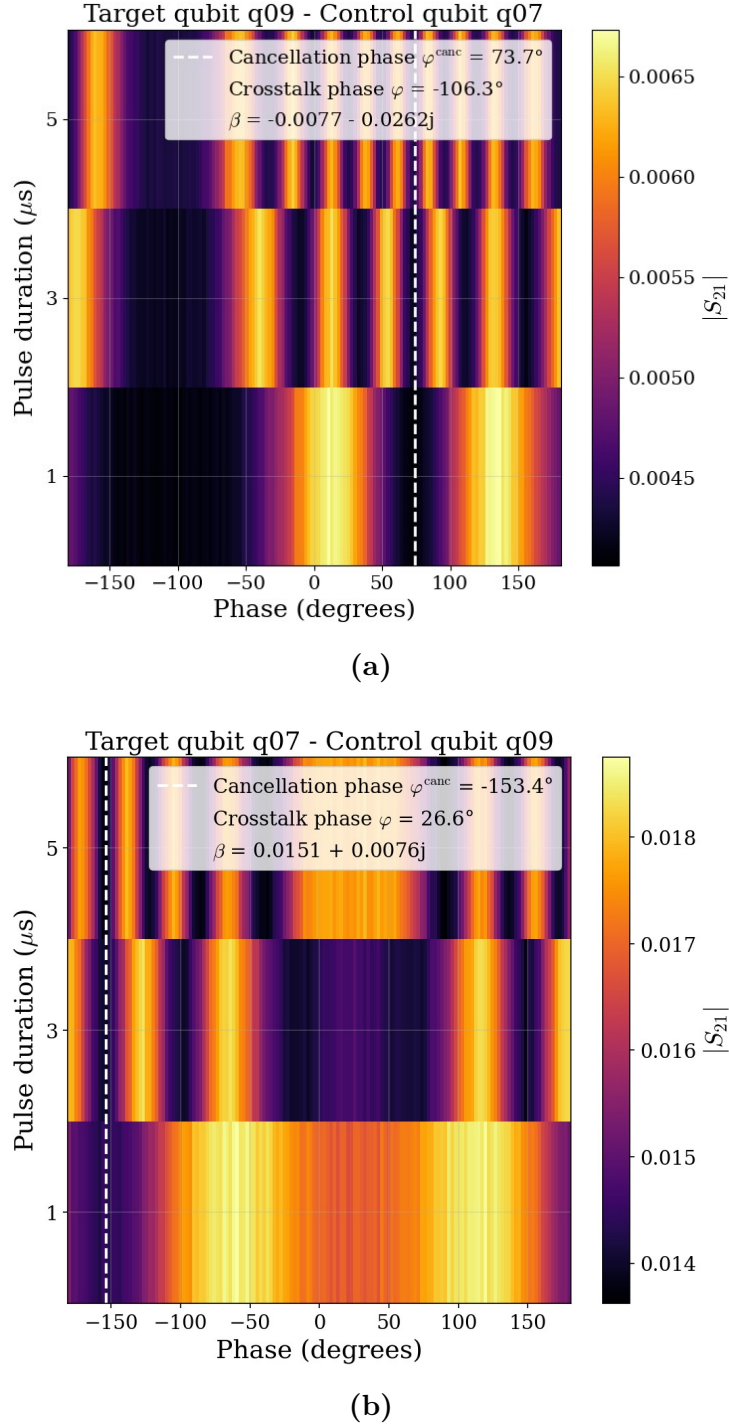


Figure 4.3: Representative $f_{01} - f_{12}$ crosstalk phase results for qubits q07 and q09. (a) Qubit q07 acts as the control qubit and q09 as the target qubit. (b) Qubit q09 acts as the control qubit and q07 as the target qubit. The white dashed line indicates the cancellation phase, at which the leaked tone and the compensation tone interfere destructively for all pulse durations. The crosstalk phase is 180° out of phase with the cancellation phase. The complex crosstalk coefficient β shown in each panel is calculated from the extracted crosstalk phase together with the amplitude $|\beta|$ obtained from the amplitude measurement (Fig. 4.2).

phase, the two tones produce the maximum net drive on the target qubit, resulting in the fastest oscillations along the pulse-duration axis. In Fig. 4.3a, this crosstalk phase point, or boundary, happens to coincide with a dark fringe for all three selected pulse durations, giving the appearance of an second cancellation point. The true cancellation phase is still the centre of the Chevron pattern, where there is no net drive on the target qubit and thus a dark fringe for any given pulse duration. It is indicated by the white dashed line in Fig. 4.3.

An example of $f_{01} - f_{01}$ crosstalk phase results obtained for qubits q11 and q15 is shown in Fig. B.2 in Appendix B.

Finally, it should be noted that, in the present setup, the qubits are driven by separate QCM-RF modules, each with its own LO (see Section 3.3). As a result, the measured cancellation phase depends on the arbitrary relative phase between the LOs of the control and target qubit. As long as both LOs remain locked, this relative phase is constant, and the calibrated crosstalk phase is stable throughout a measurement session. However, if either LO relocks, e.g., after a cluster reset, the phase-locked loop re-acquires lock at a random absolute phase, leading to a new arbitrary relative phase between the LOs. The measured crosstalk phase therefore changes and must be recalibrated. The crosstalk amplitude, however, is unaffected, as it is set solely by the hardware coupling between the drive lines and qubits. This limitation could be removed by driving both qubits from a common LO.

4.3 XY crosstalk compensation

Active XY crosstalk compensation is implemented according to the discussion in Section 2.3. An additional compensation pulse is applied to the target qubit's XY line at the control qubit's first transition frequency. This frequency is always f_{01} regardless of whether the crosstalk is of the $f_{01} - f_{01}$ or $f_{01} - f_{12}$ type, since the leaked tone always comes from the control qubit's $|0\rangle - |1\rangle$ transition. The compensation pulse has an amplitude equal to that of the intended control pulse scaled by the magnitude of the bare crosstalk coefficient, $|\beta|$, and a phase shift of 180° relative to the crosstalk signal. For the example pair q07 - q09, compensation is applied only for the crosstalk from q07 to q09, since the reverse direction is strongly suppressed by frequency detuning.

SSRO RB (see Section A.3.4 for more details) is used to evaluate the effectiveness of the compensation scheme. As can be seen in Figs. 4.5, 4.6b and 4.7b, crosstalk compensation restores the performance of qubit q09 to essentially the same level under isolated and simultaneous operations. During simultaneous RB, the average single-qubit gate error is reduced from $0.12 \pm 0.01\%$ without compensation to $0.072 \pm 0.005\%$ with compensation, in good agreement with the value of $0.069 \pm 0.006\%$ obtained for isolated RB. Likewise, the leakage rate decreases from $0.025 \pm 0.002\%$ to $0.0040 \pm 0.0008\%$, comparable to the isolated value of $0.0034 \pm 0.0001\%$. This improvement is also clearly visible from the green traces in the RB plots (Figs. 4.4 and 4.5), which represent the $|2\rangle$ state population. The performance of qubit q07 remains unchanged as expected, since no compensation pulse is applied to that qubit. The complete RB results for qubit q07, with and without crosstalk compensation, are provided in Appendix B.

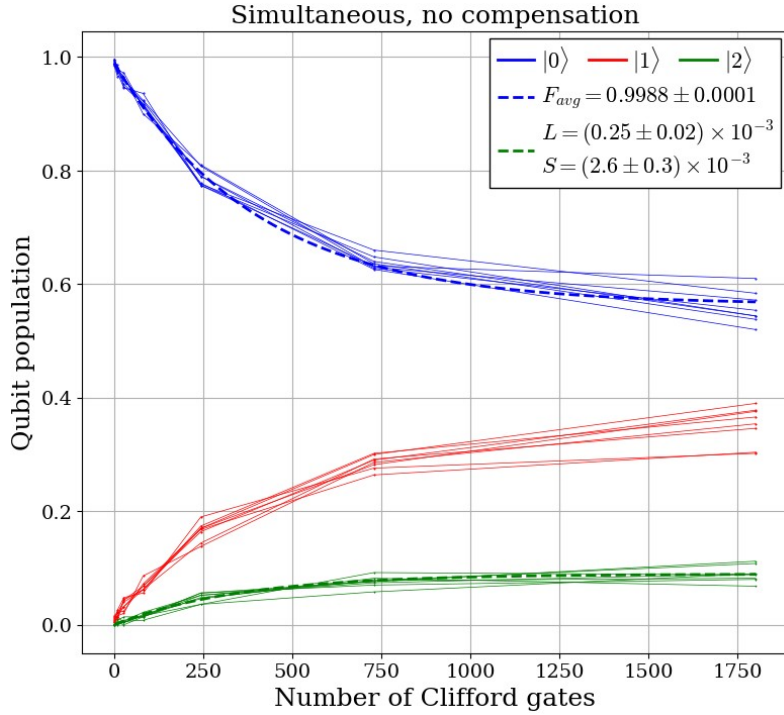
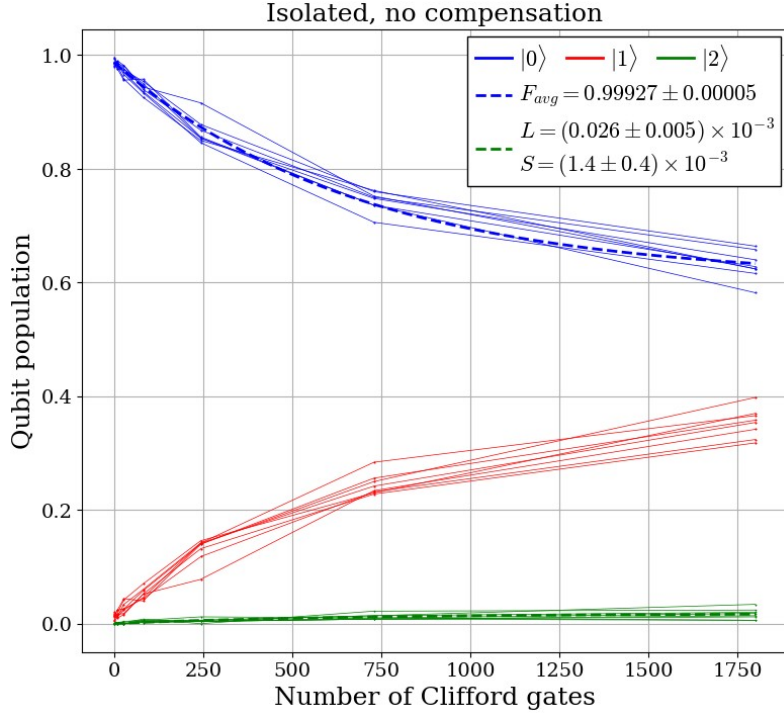


Figure 4.4: RB results for qubit q09 without crosstalk compensation. (a) Isolated RB. (b) Simultaneous RB performed with qubit q07: the $|0\rangle$ state population (blue curve) decays faster with increasing number of Clifford gates compared to the isolated case. Additionally, the $|2\rangle$ state population (green curve) is significantly higher compared to the isolated case, indicating the presence of leakage due to $f_{01} - f_{12}$ crosstalk with qubit q07.

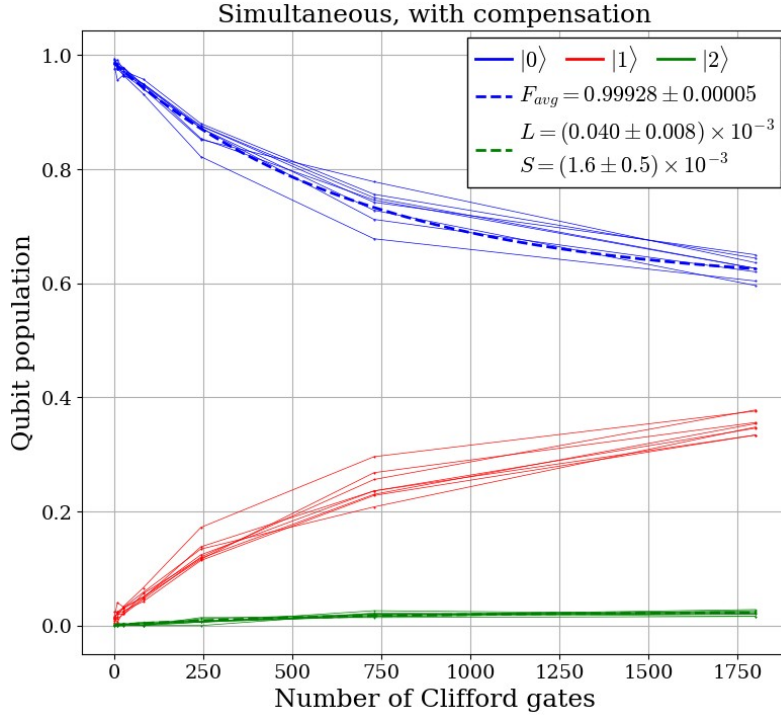
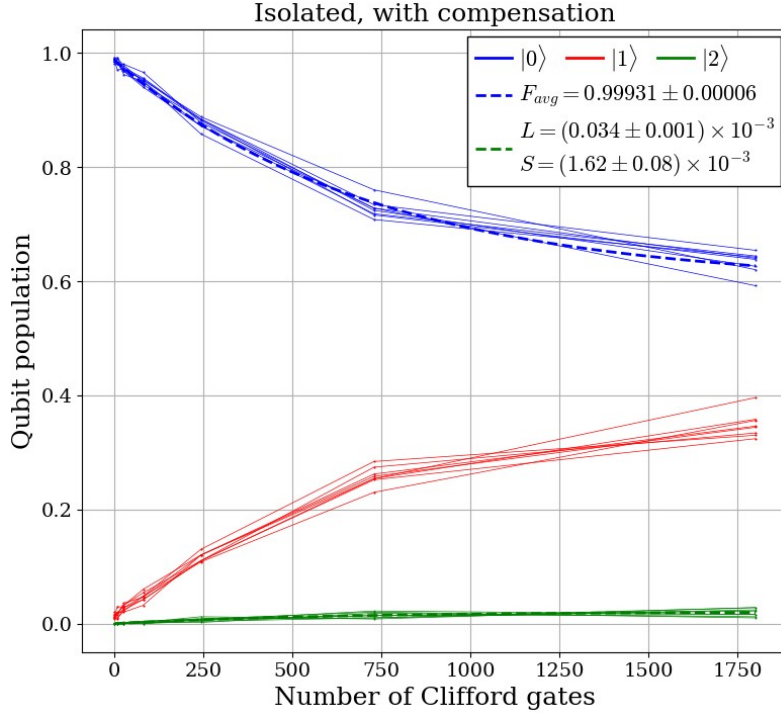


Figure 4.5: RB results for qubit q09 with crosstalk compensation. (a) Isolated RB. (b) Simultaneous RB performed with qubit q07: all curves appear essentially identical to the isolated case, suggesting that crosstalk is fully compensated and thus no leakage is present.

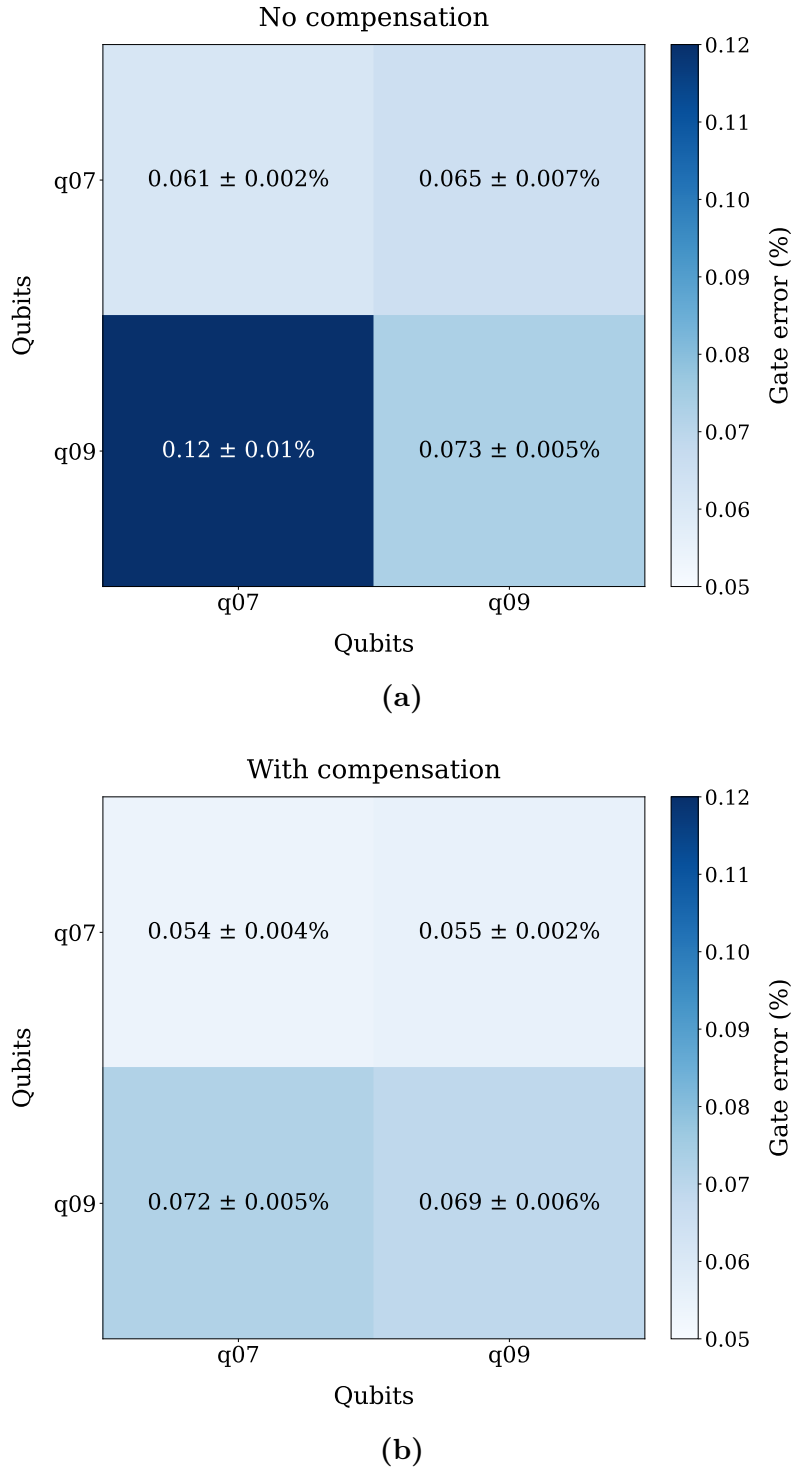


Figure 4.6: Average single-qubit gate error matrix for qubits q07 and q09. The diagonal elements correspond to isolated RB measurements, while the off-diagonal elements correspond to simultaneous RB performed on both qubits. (a) Without crosstalk compensation, qubit q09’s gate error increases by approximately 1.6 times during simultaneous operation. (b) With crosstalk compensation, qubit q09’s gate errors obtained from isolated and simultaneous RB agree within experimental uncertainties.

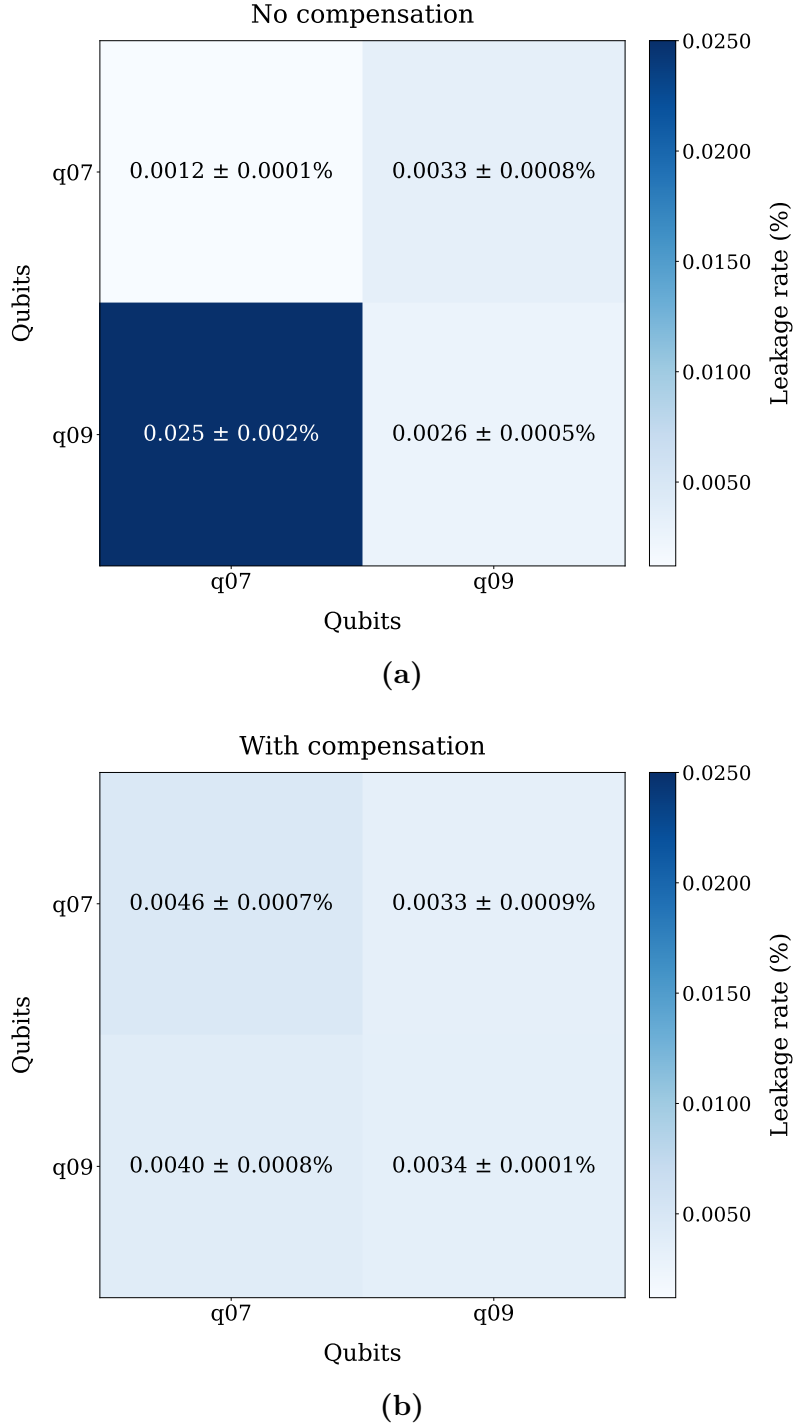


Figure 4.7: Leakage rate matrix for qubits q07 and q09. The diagonal elements correspond to isolated RB measurements, while the off-diagonal elements correspond to simultaneous RB performed on both qubits. (a) Without crosstalk compensation, qubit q09’s leakage rate increases by almost an order of magnitude during simultaneous RB. (b) With crosstalk compensation, qubit q09’s leakage rates obtained from isolated and simultaneous RB are consistent within experimental uncertainties.

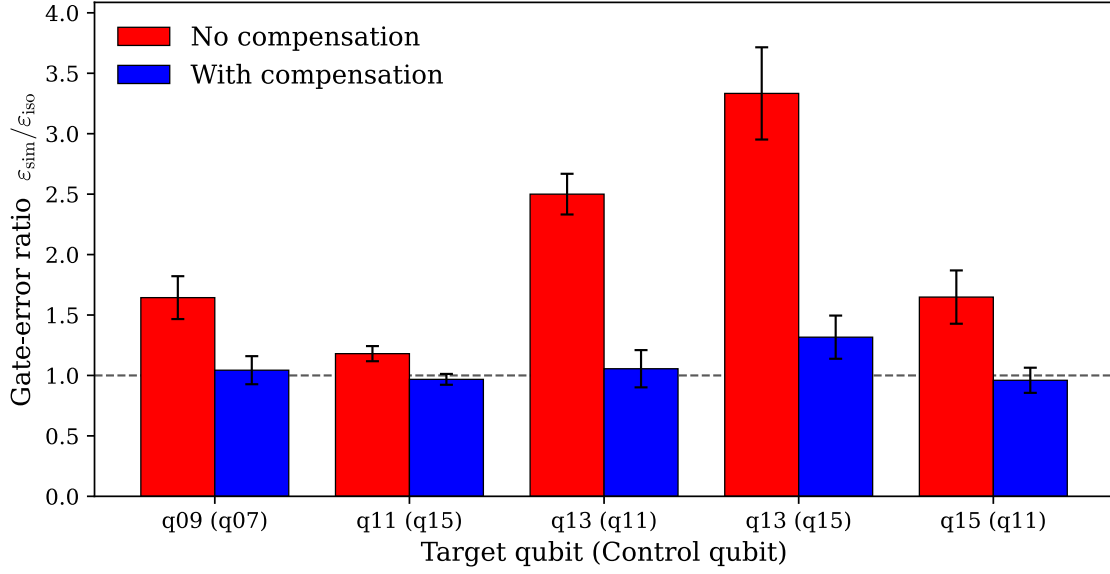


Figure 4.8: Comparison of gate-error ratios with and without XY crosstalk compensation. The gate-error ratio $\varepsilon_{\text{sim}}/\varepsilon_{\text{iso}}$ compares the average Clifford gate error obtained from simultaneous RB with that from isolated RB for each target qubit (control qubit) pair. The red bars correspond to measurements without compensation, while the blue bars correspond to those after applying the compensation scheme. The grey dashed line at $\varepsilon_{\text{sim}}/\varepsilon_{\text{iso}} = 1$ represents the ideal case, where simultaneous operations introduce no additional gate error. For most qubit pairs, the compensation reduces the ratio to approximately unity, demonstrating effective suppression of XY crosstalk.

Similar improvements in single-qubit gate fidelity are observed for three additional qubit pairs when the compensation scheme is applied: q11 – q15 ($f_{01} - f_{01}$), q11 – q13 ($f_{01} - f_{12}$), and q15 – q13 ($f_{01} - f_{12}$). Their corresponding gate-error matrices are presented in Appendix B. Fig. 4.8 compares the gate-error ratio, $\varepsilon_{\text{sim}}/\varepsilon_{\text{iso}}$, for different qubit pairs with and without crosstalk compensation, where ε_{iso} and ε_{sim} denote the average Clifford gate errors obtained from isolated RB and simultaneous RB, respectively. Without compensation, the gate error increases substantially, reaching more than two to three times its isolated value for some qubit pairs. After applying the compensation scheme, the gate-error ratio is reduced to approximately unity for almost all pairs, indicating that the effect of the crosstalk signal has been effectively eliminated.

5

Conclusion & Future Outlook

5.1 Conclusion

This thesis has successfully developed and validated an experimental protocol for measuring and compensating XY crosstalk on Chalmers' 25-qubit fixed-frequency transmon quantum processor. Separate experiments are implemented to extract the amplitude and phase of the complex crosstalk coefficient for both crosstalk channels, $f_{01} - f_{01}$ and $f_{01} - f_{12}$. The measured coefficients are integrated into a crosstalk matrix and used for active compensation at the signal level, which is incorporated into the existing automated calibration software framework.

The effectiveness of the compensation scheme is demonstrated through isolated and simultaneous SSRO RB. For selected qubit pairs, the average single-qubit gate error during simultaneous operation is reduced by up to 2.5 times to the same level as the isolated single-qubit gate error. For pairs exhibiting $f_{01} - f_{12}$ crosstalk, a substantial reduction in leakage rate (e.g., by almost 10 times for qubit pair q07 - q09) is also observed. These results demonstrate the effective characterisation and mitigation of XY crosstalk through active signal compensation.

This work has contributed to a more comprehensive understanding of XY crosstalk in Chalmers' 25-qubit quantum processor and improved its gate performance, enabling high-fidelity parallel operations. The developed methods are not limited to the current device architecture and can also be adapted to other fixed-frequency superconducting quantum processors. Furthermore, this thesis has established a clear distinction between the different definitions of crosstalk coefficients (β , β^{eff} , and $\tilde{\beta}$), providing a consistent framework for their interpretation and application when dealing with XY crosstalk.

Although only demonstrated for selected qubit pairs within one cluster of the processor, the developed methodology is inherently scalable and provides a practical pathway towards automated XY crosstalk characterisation and compensation in larger systems.

5.2 Future outlook

Several directions can be pursued to further extend the work presented in this thesis:

- Extend the XY crosstalk characterisation and compensation to the remaining qubit pairs in both clusters A and B, and ultimately to simultaneous compensation of all significant crosstalk pathways across the processor. This

will require addressing several practical constraints as the number of applied compensation tones increases, such as the number of available sequencers per QCM-RF module and the total waveform memory limitation.

- Improve qubit calibration to further enhance the reliability of crosstalk characterisation and compensation.
- Investigate the use of the effective crosstalk coefficient β^{eff} as a criterion for pruning the crosstalk matrix and identifying qubit pairs that require compensation. This can be achieved by experimentally determining β^{eff} for different qubit pairs and studying its correlation with the gate-error ratio $\varepsilon_{\text{sim}}/\varepsilon_{\text{iso}}$ in the absence of compensation. Such an analysis would reveal whether a threshold value of β^{eff} exists below which crosstalk has a negligible impact on gate performance.
- Perform a study of the dependence of XY crosstalk on qubit separation. As the processor grows, the number of frequency collisions will increase. Understanding how qubit separation affects the strength of crosstalk is important for determining how crosstalk mitigation should scale with the processor size.
- Ultimately, a more comprehensive study separating XY crosstalk contributions from different physical mechanisms would provide guidance for future device designs and enable crosstalk to be mitigated also at the hardware level, rather than solely relying on active compensation schemes at the software level.

Appendix A

Single-Qubit Gate Calibration & Characterisation

This appendix describes the calibration of single-qubit gates, i.e., the sequence of experiments used to optimise SSRO and the microwave pulses used for single-qubit operations. The ultimate goal is to maximise both the single-qubit gate fidelity and SSRO fidelity.

To achieve this, many parameters of the quantum processor, such as the resonator frequencies, qubit frequencies, and pulse amplitudes, need to be measured and optimised. This is done by performing a series of experiments, each of which is designed specifically to calibrate one of these quantities. A clear dependency exists between these parameters. For instance, the readout frequency and amplitude must be calibrated before qubit measurements can be performed reliably. These dependencies are naturally represented as a graph, where each experiment corresponds to a node. The dependency graph for the single-qubit calibration chain is shown in Figure A.1.

The first four sections of this appendix explain the basics of individual nodes and present representative results obtained for qubits in cluster A. Rather than following the exact order of the dependency graph, the nodes are grouped according to their functionality:

- Readout calibration nodes (Section A.1) optimise qubit measurement.
- Qubit calibration nodes (Section A.2) determine and optimise the pulse parameters required for single-qubit control.
- Qubit characterisation nodes (Section A.3) include coherence time measurements and RB.
- Crosstalk measurement nodes (Section A.4) are the ones developed and described in this thesis.

More information on most of these nodes can be found in Refs. [76–78].

The appendix concludes with Section A.5, which presents an example of a complete single-qubit gate bring-up procedure. The dependency graph in Fig. A.1 represents only the minimum set of experiments required to reach a given node. In practice, however, a full calibration requires a higher level of complexity as some experiments need to be repeated multiple times with refined sweep range or updated settings to extract the relevant parameters with high accuracy. This is particularly relevant for the early steps of the calibration chain, such as resonator and qubit

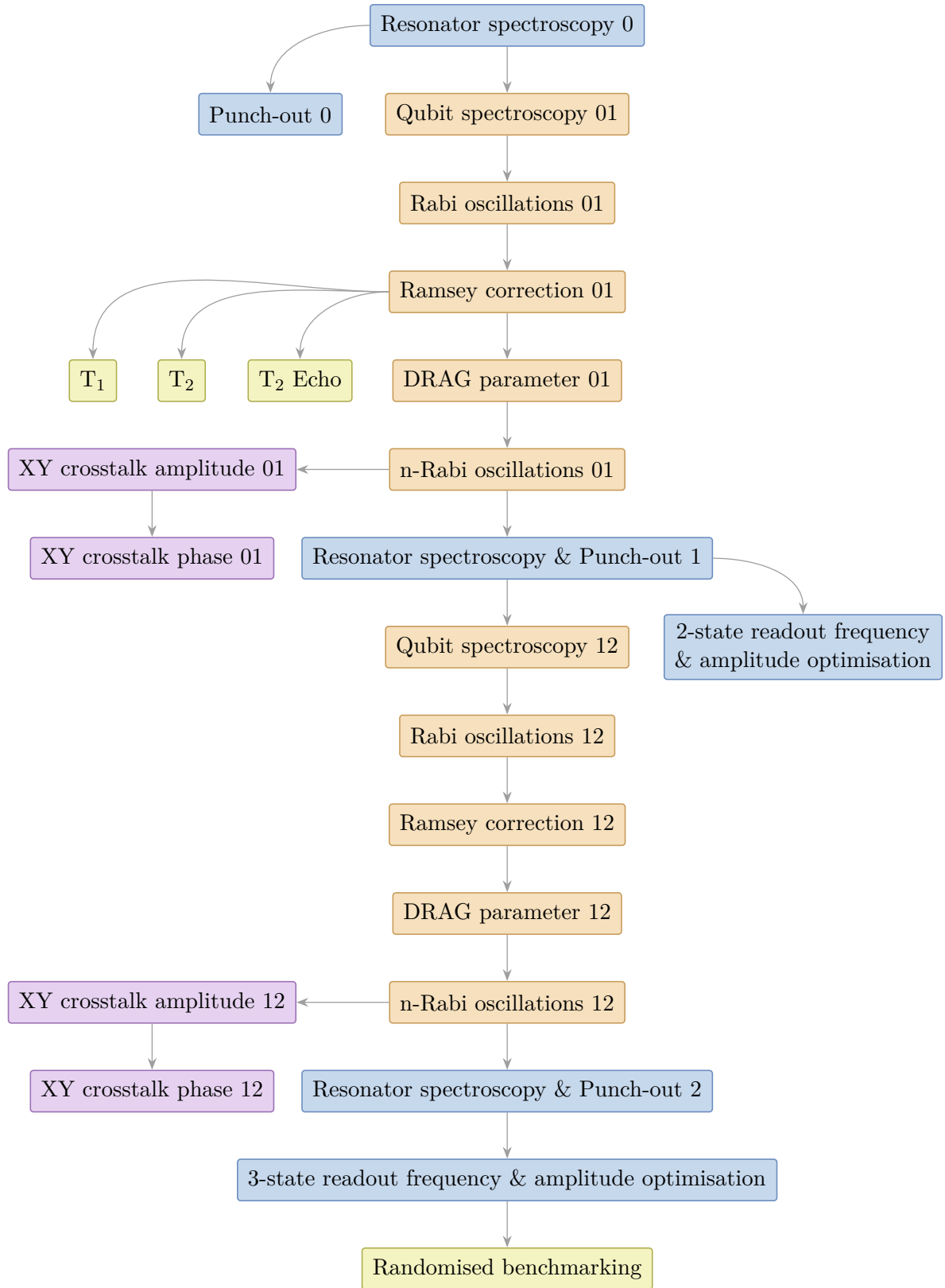


Figure A.1: Single-qubit gate calibration node dependency tree. Blue: readout calibration nodes, orange: qubit calibration nodes, yellow: qubit characterisation nodes, purple: crosstalk measurement nodes.

spectroscopy, where a coarse, wide frequency sweep is first used to find the resonance, followed by one or more finer sweeps over a narrower range to determine the frequency with higher precision.

A.1 Readout calibration nodes

A.1.1 Resonator spectroscopy

A single-qubit gate calibration sequence typically starts with determining the frequency of the readout resonator. This is achieved by sweeping a probe tone around the expected resonator frequency while measuring the transmission coefficient $|S_{21}|$. The resonator frequency corresponds to the minimum of the Lorentzian dip in the transmission spectrum (Fig. A.2). As discussed in Section 1.3.4, the resonator frequency depends on the qubit state, with the separation between adjacent resonator frequencies given by the dispersive shift 2χ . The measurement schedules for resonator spectroscopy 0, 1, and 2 are almost identical; the only difference is in the qubit initialisation step.

As can be seen in Fig. A.2, some resonator responses (e.g., qubits q14 and q12) are distorted from an ideal Lorentzian shape and appear asymmetric. Optimising the readout pulse amplitude is thus required and is the purpose of punch-out, which is described in the following section.

A.1.2 Punch-out

Punch-out is the short name for resonator spectroscopy as a function of probe amplitude, and is used to optimise the readout amplitude. The physical motivation behind this experiment is the requirement for the dispersive approximation to be valid, where the average photon number n in the resonator has to satisfy $n \ll \Delta^2/g^2$ (Δ is the frequency detuning and g is the qubit - resonator coupling strength). Under this condition, the qubit and the resonator are dispersively coupled and the resonator exhibits a stable qubit-state-dependent dressed frequency. As the probe amplitude is increased, the photon number in the resonator also increases and eventually drives the system out of the dispersive regime. The resonator frequency would then shift towards its bare frequency, which is lower than the dressed frequency. At the same time, the readout power should be as high as possible, since stronger readout signals lead to faster and more accurate state discrimination. Thus, by tracking the resonator frequency as a function of probe amplitude, it is possible to estimate the optimal readout amplitude, which is the highest possible amplitude before the resonator *punches out*.

The frequency sweep is centred on the resonator frequency obtained from a preceding resonator spectroscopy experiment. For each probe amplitude, the resonator response is fitted to extract the resonator frequency. The frequencies obtained from each fit are then compared, and if the difference between two consecutive points exceeds a predefined threshold (e.g., 50 kHz), the corresponding amplitude is marked as the punch-out endpoint (the red cross in Fig. A.3). The resonator is also considered to have punched out if the goodness-of-fit metric falls below a predefined threshold

for more than three consecutive amplitudes. This criterion is useful because the fit quality worsens in the transition region between the dispersive and punch-out regimes, so it can provide an additional indicator of the punch-out endpoint. The optimal readout amplitude (the red line in Fig. A.3) is chosen as 80 % of the endpoint amplitude to provide some safety margin from the onset of punching-out.

It should be noted that punch-out is used only to extract the optimal readout amplitude, not the resonator frequency. After the optimal amplitude has been identified, resonator spectroscopy is typically performed once more using this value and a narrower frequency range to fine-tune the resonator frequency. An example is shown in Fig. A.4, where the resonator responses exhibit a more symmetric Lorentzian shape after punch-out.

As with resonator spectroscopy, punch-out nodes are performed for all relevant qubit states - $|0\rangle$, $|1\rangle$, and $|2\rangle$, differing only in the state preparation step.

A.1.3 State-discrimination readout frequency optimisation

The goal of this experiment is to find an optimal single frequency value for two-state and three-state dispersive readout, i.e., one that best separates the qubit states in the IQ plane, as described in Section 1.3.4. The qubit is first prepared in each of its two (or three) lowest eigenstates, after which square readout pulses are applied at different frequencies. For each frequency point, the returned signals for each state are integrated into single complex points in the IQ plane. The amplitude of the square probe pulse is the same for all states and is chosen as the smallest of the readout amplitudes obtained previously from punch-out experiments of the corresponding states.

For two-state readout, the complex IQ distance between the average points of states $|0\rangle$ and $|1\rangle$ is calculated at each frequency. The optimal readout frequency is then the one that maximises this distance, where the two states are most clearly distinguishable. In the case of three-state readout, the pairwise IQ distances between the average responses of the three states, d_{01} , d_{12} , and d_{02} , are first computed. These distances define a triangle in the IQ plane, whose area is calculated using Heron's formula:

$$A = \sqrt{s(s - d_{01})(s - d_{12})(s - d_{02})}, \quad (\text{A.1})$$

where s is the semi-perimeter of the triangle:

$$s = \frac{d_{01} + d_{12} + d_{02}}{2}. \quad (\text{A.2})$$

The optimal readout frequency is then the one that maximises the triangle area A , corresponding to the case where all three states are mutually well separated. As expected, this frequency is typically located between the state-dependent resonator frequencies, as observed in Fig. A.5a.

A.1.4 State-discrimination readout amplitude optimisation

While the node described in the previous section finds the optimal readout frequency, this node determines the optimal readout amplitude. As before, the qubit is prepared in each of its relevant states and probed with a square readout pulse.

The frequency of the pulse is fixed to the optimised value obtained from the previous node, while the pulse amplitude is swept over a range determined from the punch-out amplitudes. Unlike the frequency optimisation experiment, where only the average IQ responses are considered, in this experiment, all single shots are kept. For each amplitude value, a linear discriminant analysis classifier is trained to find the linear decision boundaries that best separate the state clusters in the IQ plane. A confusion matrix is then constructed, from which the assignment fidelity, defined as the average probability of correctly identifying the prepared state, is extracted. The optimal readout amplitude is the one maximising this assignment fidelity (Fig. A.5b).

For three-state readout, the classifier is then re-fitted using the optimal amplitude, yielding three decision boundary angles and a centroid (i.e., the intersection point of the three boundaries). These boundaries slice the IQ plane into three distinct regions, each associated with one of the states $|0\rangle$, $|1\rangle$, and $|2\rangle$, as shown in Fig. A.6. Future single-shot measurements can then be assigned to a state according to the region in which their IQ point falls. The same principle applies to two-state readout, except that only a single decision boundary is needed.

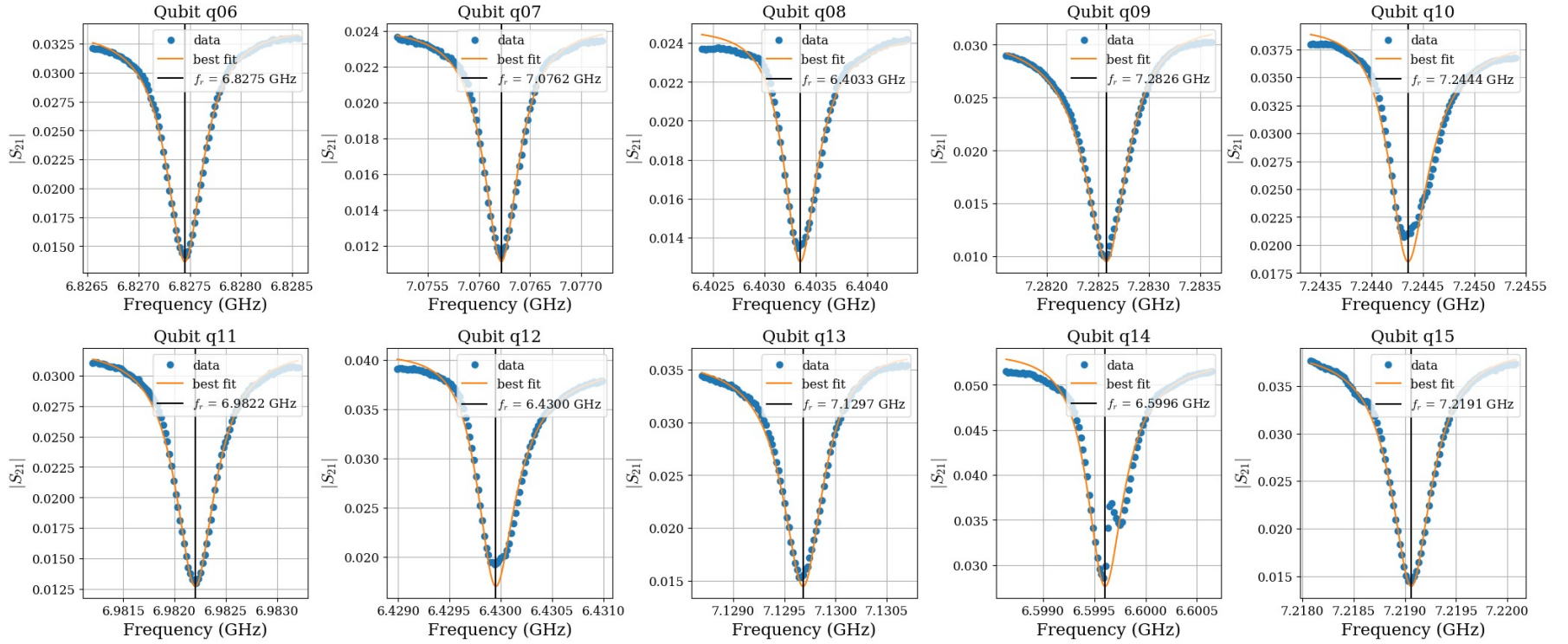


Figure A.2: Resonator spectroscopy for state $|0\rangle$. In this example, the probe pulse amplitudes were not optimised, leading to a non-Lorentzian dip for some resonators.

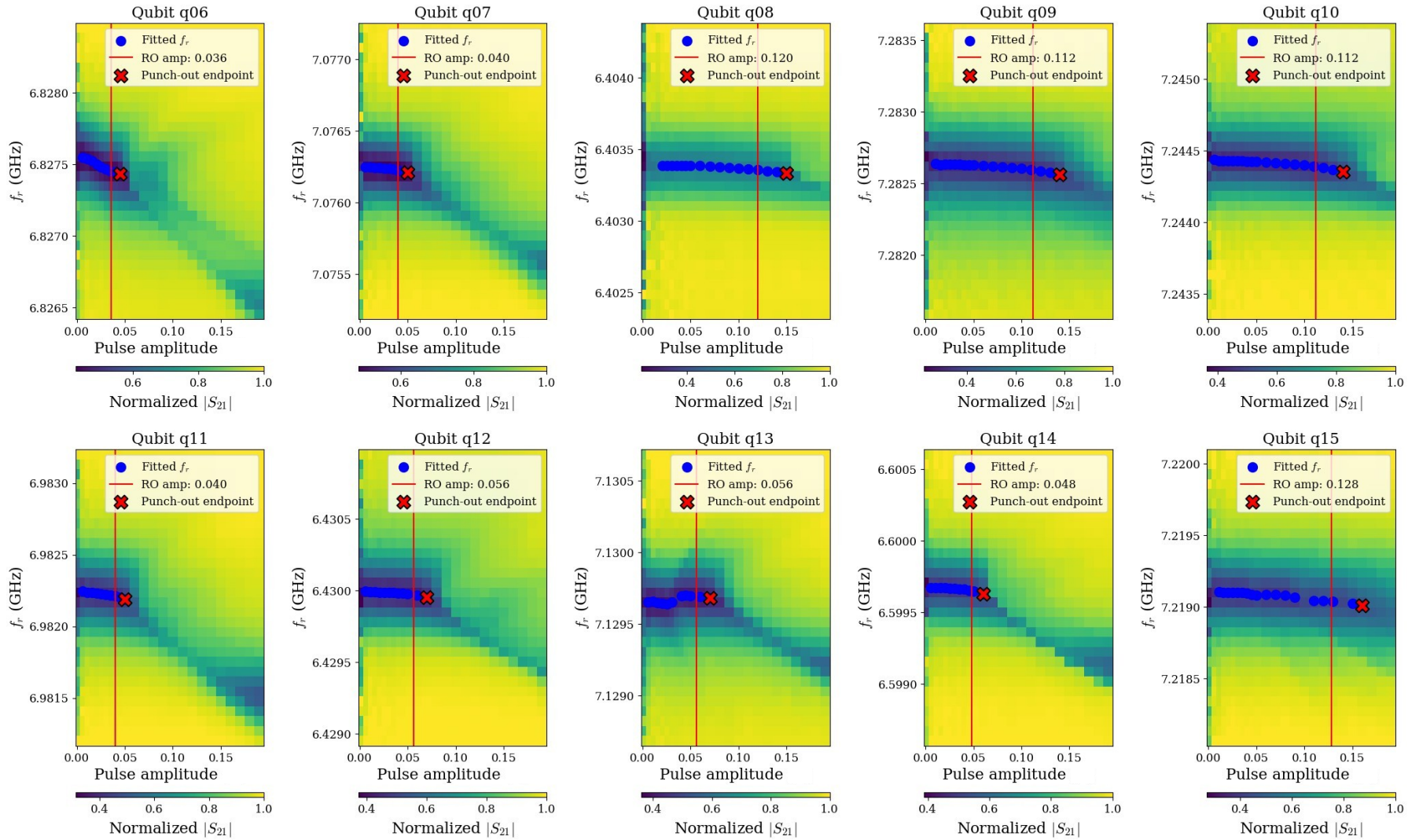


Figure A.3: Punch-out for state $|0\rangle$. The red cross denotes the punch-out endpoint, while the red line marks the optimal readout pulse amplitude.

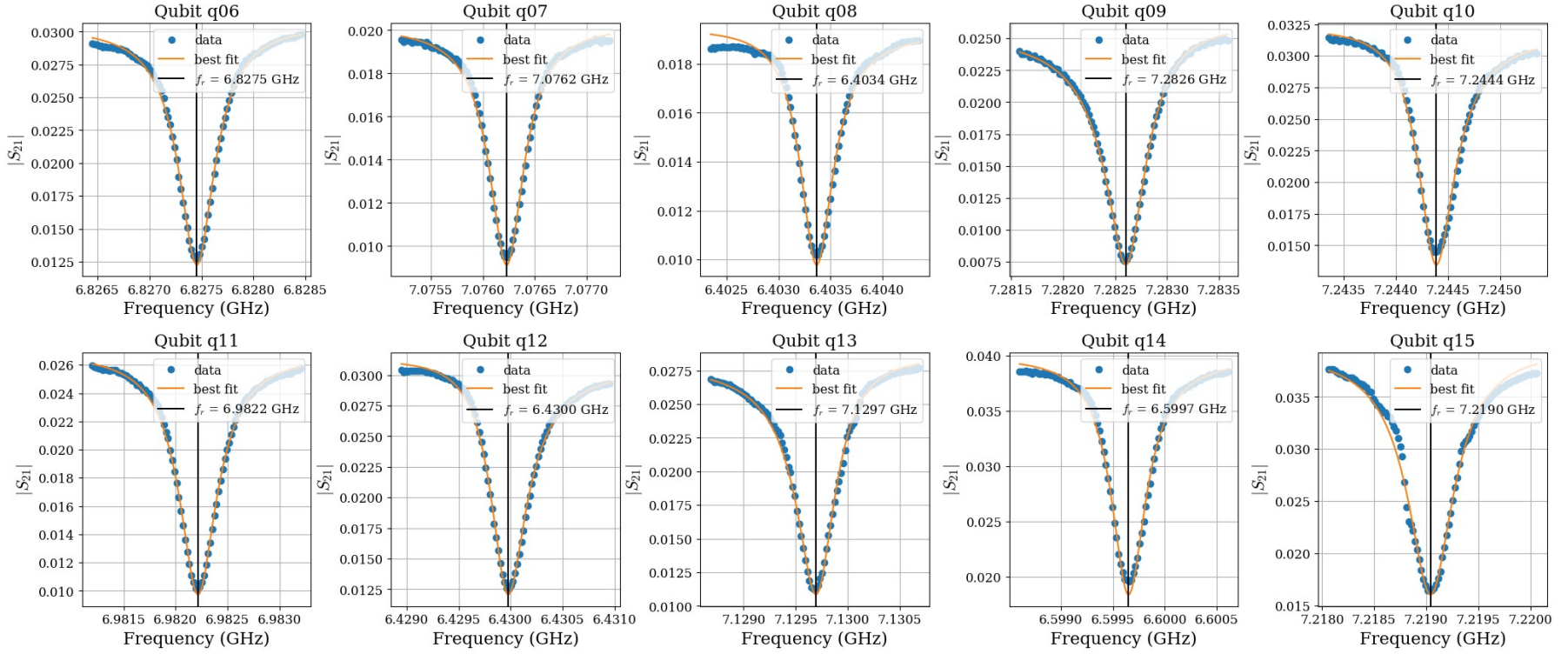


Figure A.4: Readout spectroscopy for state $|0\rangle$ after punch-out. With the optimal readout amplitude, all dips are now Lorentzian.

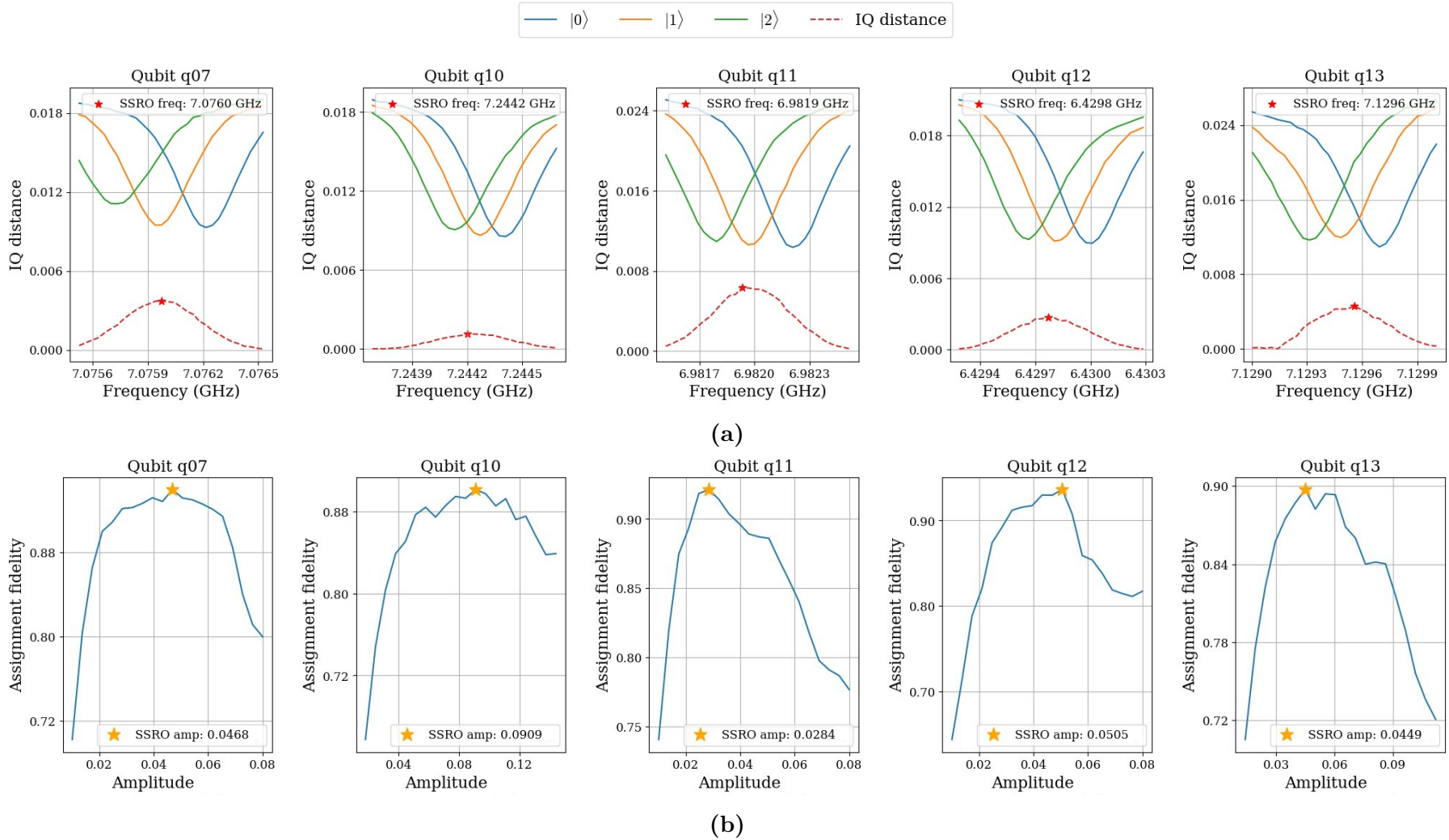


Figure A.5: Three-state readout optimisation. (a) Frequency optimisation: The red star denotes the optimal readout frequency that gives the best separation of qubit state clusters in the IQ plane. The blue, orange, and green lines are the $|S_{21}|$ traces of the three qubit states at this frequency (axes not shown). (b) Amplitude optimisation: The yellow star represents the optimal readout amplitude where the qubit state assignment fidelity is the highest.

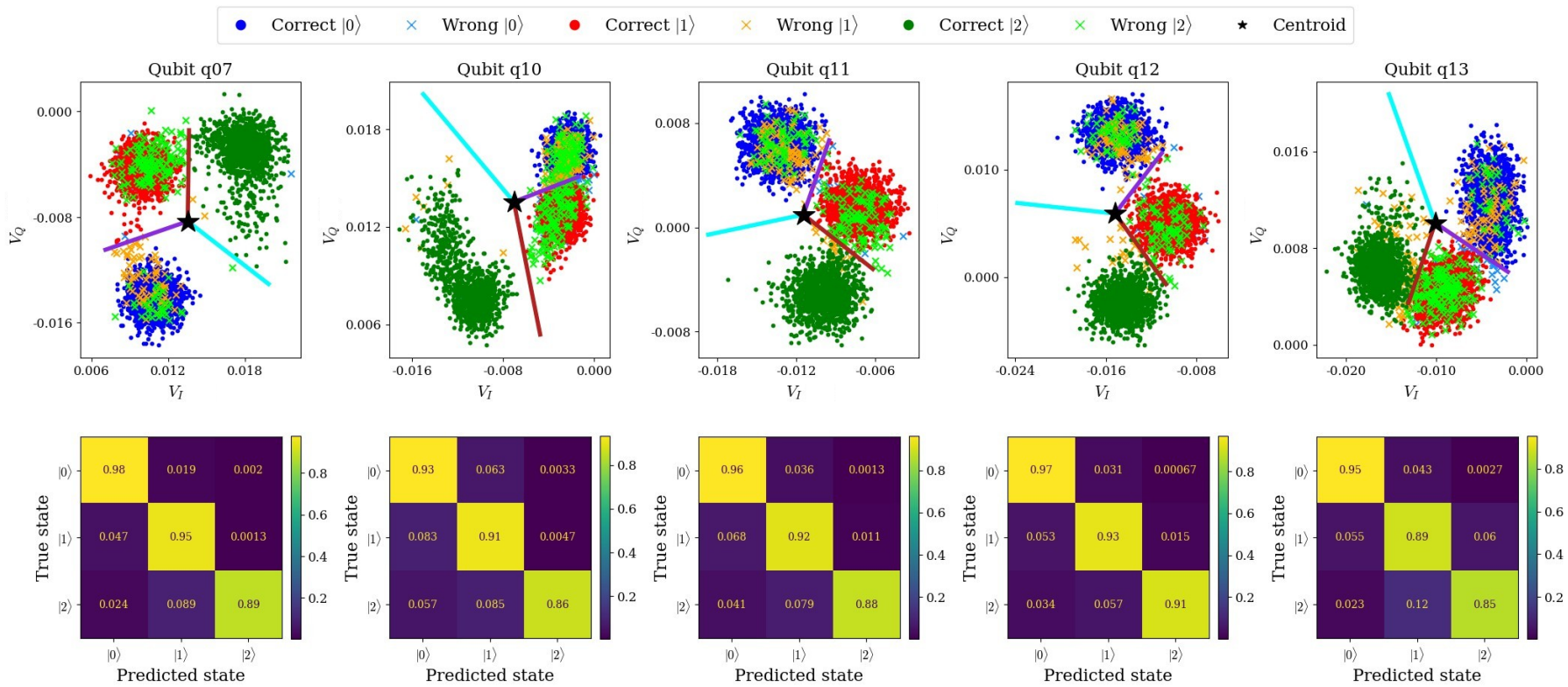


Figure A.6: Three-state SSRO with optimised readout frequency and amplitude. Top: Three decision boundaries divide each IQ plane into three areas corresponding to the three qubit states $|0\rangle$, $|1\rangle$, and $|2\rangle$. Bottom: Confusion matrices showing the rates of correctly and incorrectly identifying the prepared states.

A.2 Qubit calibration nodes

A.2.1 Qubit spectroscopy

Qubit spectroscopy is performed to determine the qubit's first and second transition frequencies (f_{01} and f_{12} , respectively). The qubit is first prepared in the appropriate initial state ($|0\rangle$ for the measurement of f_{01} and $|1\rangle$ for f_{12}) and then driven by a microwave pulse whose frequency is swept around the expected transition frequency. The qubit is subsequently measured via the readout resonator. When the drive frequency matches the qubit's transition frequency, the qubit is excited to a higher energy level, leading to a dispersive shift of the resonator frequency, and hence an increase in the resonator's transmission signal. After each measurement, the qubit is reset to its ground state before proceeding to the next frequency point. By sweeping over a range of frequencies and fitting the resulting Lorentzian peak in the transmission spectrum, the qubit frequency can be extracted.

In addition to the drive frequency, the pulse amplitude can also be swept simultaneously in a two-dimensional scan (Fig. A.7). The drive amplitude should be large enough to produce a clear signal, but not so large that it introduces power broadening or drives unwanted transitions. The point with the strongest response is identified using peak-detection algorithms, giving both the qubit frequency and drive amplitude to be used in subsequent calibration experiments [77].

A.2.2 Rabi oscillations

This experiment aims to calibrate the amplitude of the π pulse used for single-qubit control. In a Rabi oscillations 01 experiment, the qubit is initialised in its ground state and then driven at its resonance frequency, f_{01} using a Gaussian pulse of fixed duration and varying amplitude. The drive induces coherent rotations of the qubit state on the Bloch sphere about an axis in the XY plane (see Section 1.3.3), resulting in a periodic population transfer between the ground state and the first excited state. Measuring the qubit would thus yield sinusoidal oscillations of the resonator's response as a function of the drive amplitude, also known as Rabi oscillations (Fig. A.8). The optimal π pulse amplitude, A^π , corresponds to a full population transfer between states $|0\rangle$ and $|1\rangle$. In practice, this value is obtained by fitting the measured oscillations and finding the amplitude corresponding to the first maximum from the origin, or equivalently, as half of the oscillation period in amplitude space. The same principle applies to the Rabi oscillations 12 experiment, from which the π pulse amplitude used to drive the $|1\rangle - |2\rangle$ transition can be extracted.

A.2.3 Ramsey correction

This experiment is used to refine the qubit's transition frequency after an initial estimate has been obtained from qubit spectroscopy. The sequence consists of two $\pi/2$ pulses separated by a varying delay time τ . The pulses are intentionally applied with a small known detuning from the qubit frequency, which is referred to as the artificial detuning. During the delay, the qubit accumulates a relative phase at a rate determined by the detuning between the drive frequency and the actual qubit

frequency. As a result, the measured qubit population oscillates as a function of τ , producing Ramsey fringes. By fitting the oscillation frequency of these fringes, the residual detuning between the drive and the qubit frequency can be accurately extracted (fitted detuning). The qubit frequency can then be corrected by this fitted detuning (Fig. A.9). The Ramsey sequence is also used to measure the qubit's dephasing time T_2 (see Section A.3.2).

A.2.4 DRAG parameter calibration

The goal of this experiment is to determine the optimal value of the DRAG parameter β_D for the DRAG pulse (see Section 2.1.2). The qubit is first initialised in the appropriate state ($|0\rangle$ or $|1\rangle$). A sequence consisting of an X gate followed by its inverse is applied multiple times before the qubit is measured. If β_D is not properly calibrated, small coherent errors will accumulate throughout the pulse sequence, producing oscillations in the measured population as a function of β_D . The optimal β_D value is the one that minimises the accumulated error, which also corresponds to the best suppression of leakage (Fig. A.10).

In the Quantify Scheduler framework, the DRAG pulse is defined slightly differently from the form introduced in Section 2.1.2 as:

$$\Omega(t) = \Omega_0 e^{-t^2/2\sigma^2} - i\gamma \frac{t}{\sigma} e^{-t^2/2\sigma^2}, \quad (\text{A.3})$$

where γ is the unitless amplitude of the derivative component. Comparison with Eq. 2.18 gives:

$$\gamma = \frac{\beta_D}{\sigma}. \quad (\text{A.4})$$

The parameter that is optimised by this experiment, plotted in Fig. A.10, and supplied directly to the software is therefore γ , rather than β_D .

A.2.5 n-Rabi oscillations

The purpose of this experiment is to fine-tune the π pulse amplitude obtained from the Rabi oscillations experiment (Section A.2.2). The principle is similar to the DRAG calibration experiment: a sequence of two consecutive X gates is applied to the qubit multiple times before it is read out. Ideally, each pair of X gates corresponds to a 2π rotation and should bring the qubit back exactly to its initial state. However, if the pulse amplitude deviates slightly from its correct value, small coherent errors will be present and become increasingly visible as the number of X gate pairs increases. By sweeping the π pulse amplitude around the value obtained from Rabi oscillations and fitting the resulting oscillation patterns, it is possible to extract a more accurate amplitude value for which the error is consistently minimal across all repetition numbers (Fig. A.11).

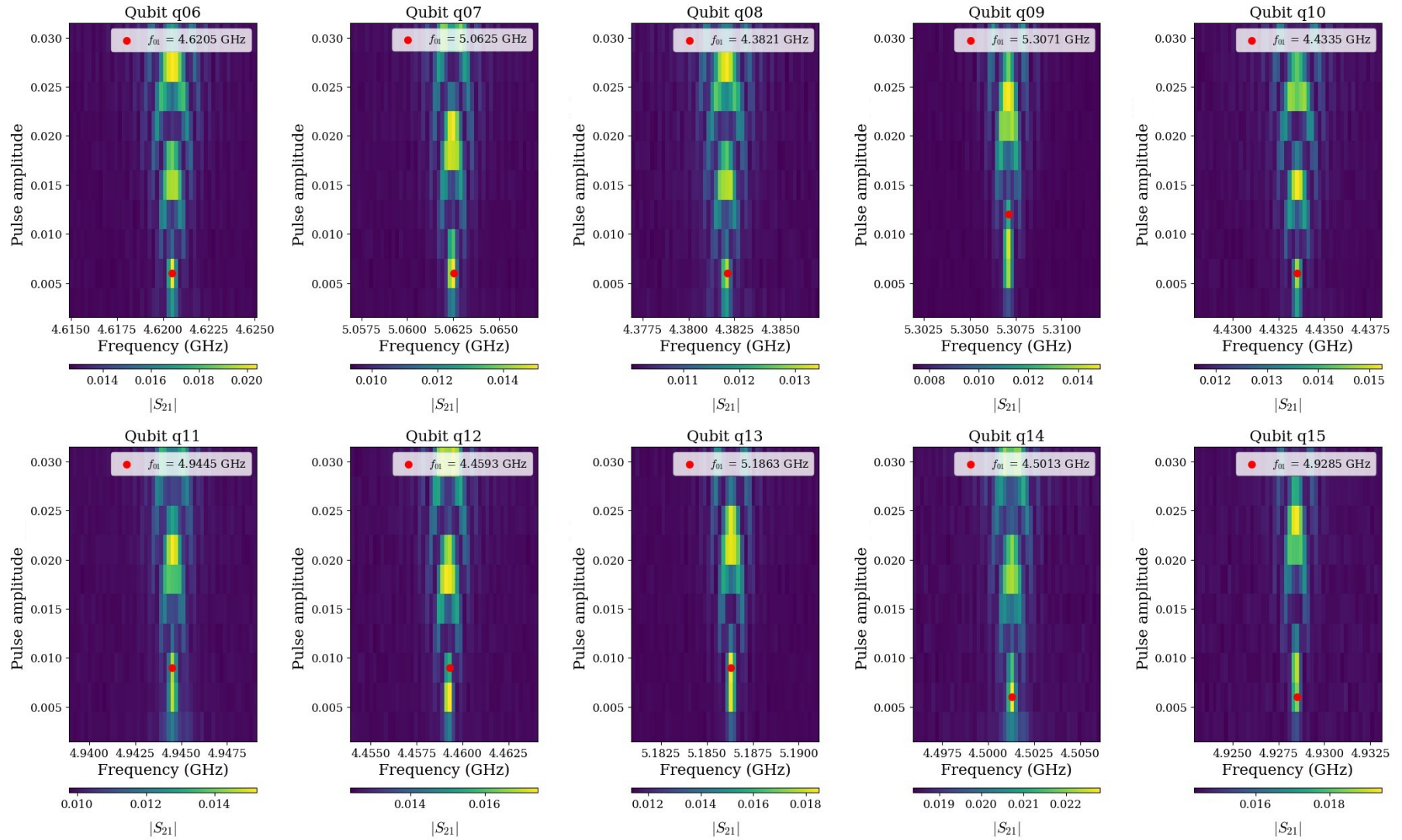


Figure A.7: Qubit spectroscopy for the $|0\rangle - |1\rangle$ transition. The red dot marks both the qubit's resonance frequency and the optimal drive amplitude.

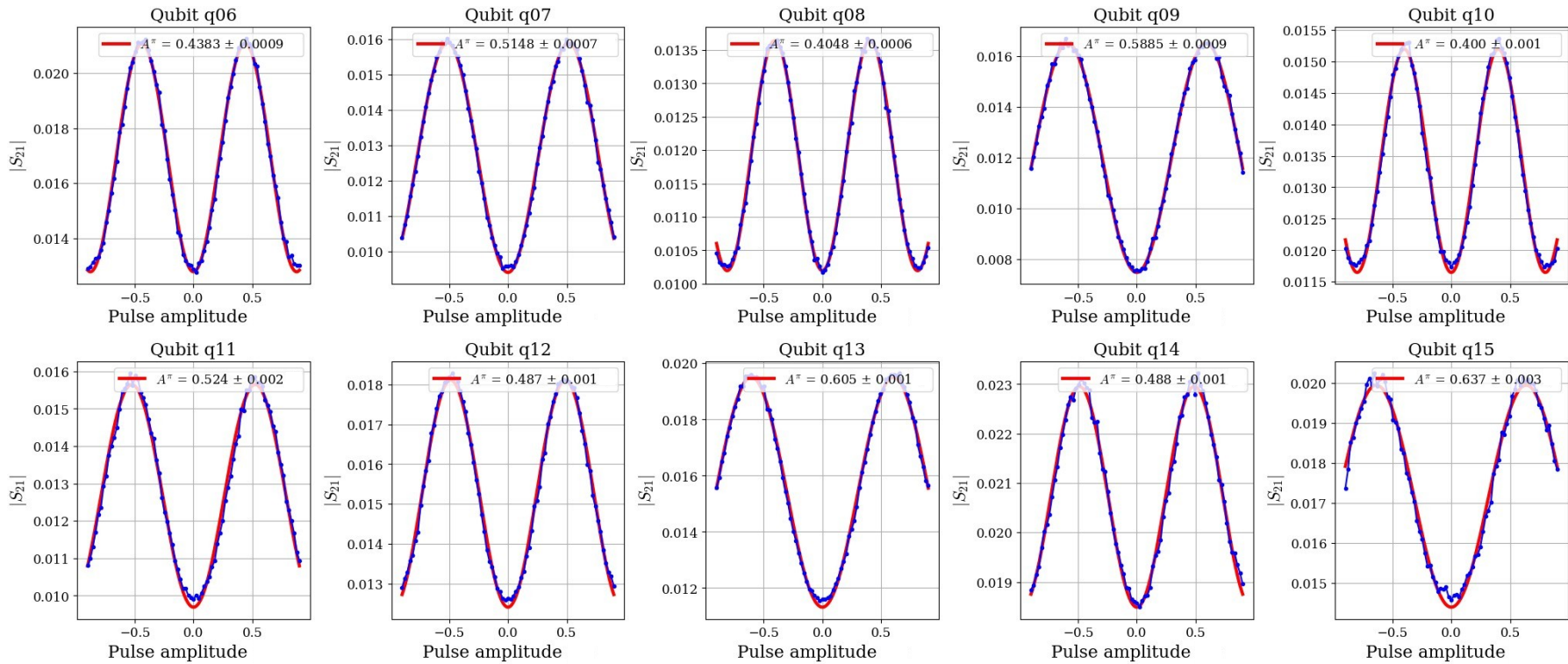


Figure A.8: Amplitude-swept Rabi oscillations for the $|0\rangle - |1\rangle$ transition. The π pulse amplitude A^π is extracted as the amplitude corresponding to the first maximum from the origin, or as half the peak-to-peak distance of the Rabi oscillations.

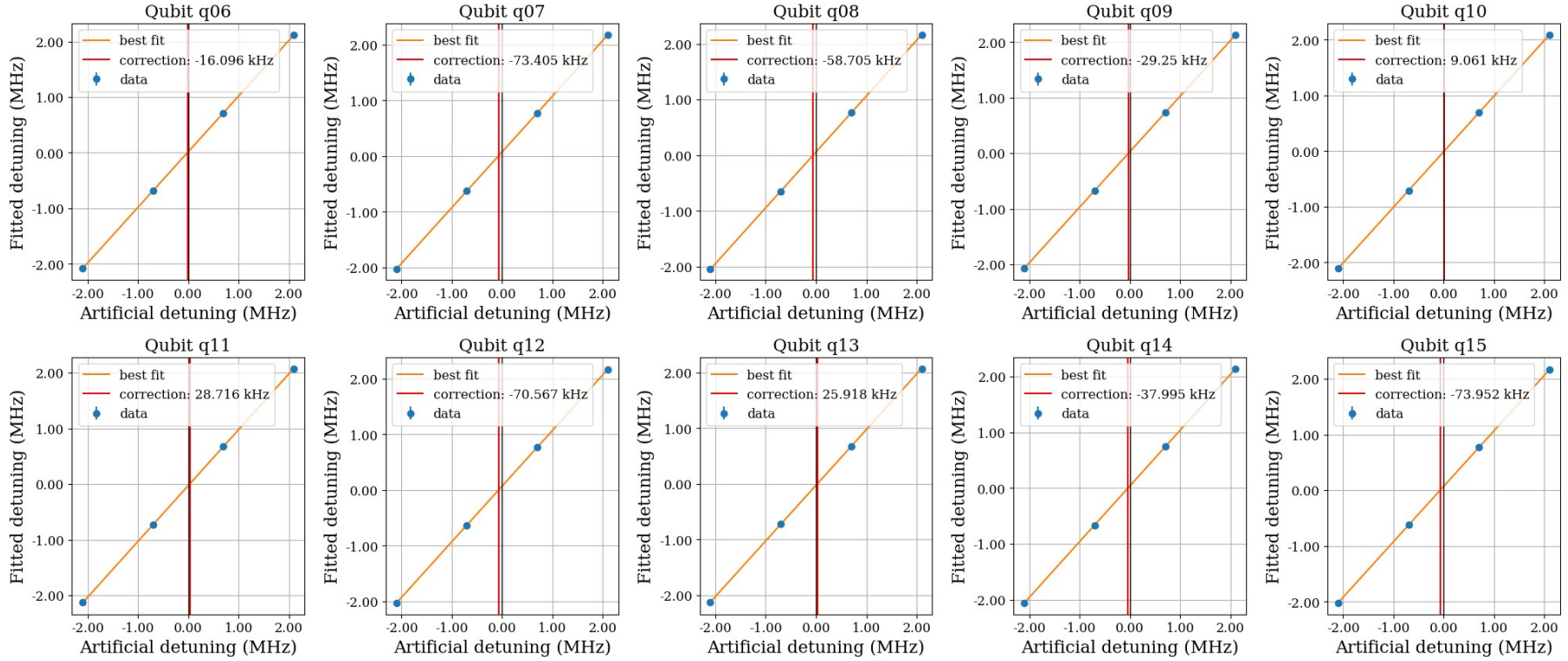


Figure A.9: Ramsey correction for the $|0\rangle - |1\rangle$ transition. The black line indicates zero artificial (applied) detuning, while the red line goes through the zero fitted detuning point. The horizontal offset between these two lines reflects how much the qubit frequency needs to be corrected. In this example, the corrections were small for all qubits, only on the order of tens of kHz, indicating that the qubit frequencies had been well calibrated from qubit spectroscopy.

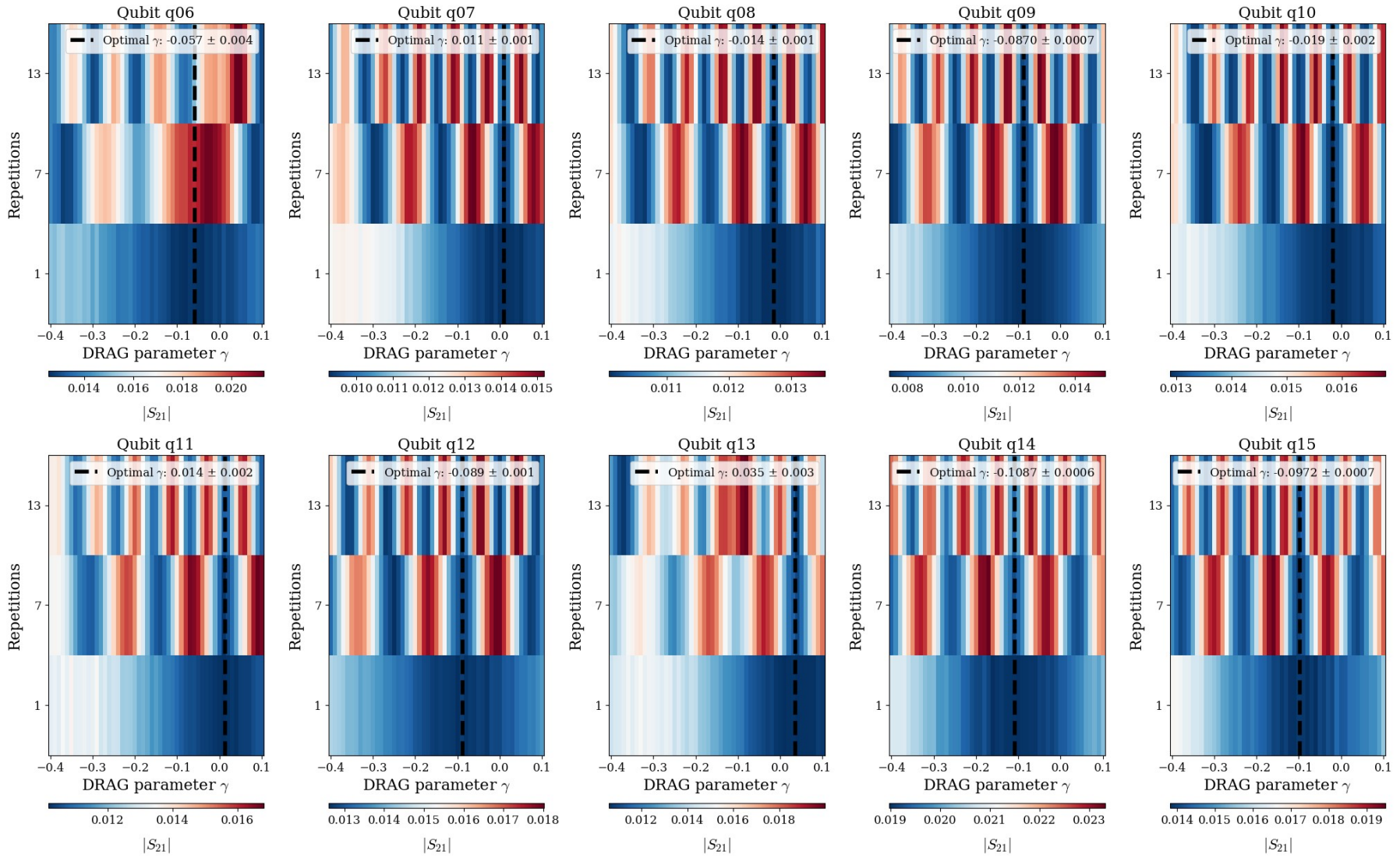


Figure A.10: DRAG parameter optimisation for the $|0\rangle - |1\rangle$ transition. The black dashed line indicates the optimal γ value, where a minimum in $|S_{21}|$ is found across different numbers of X gate repetitions. In this example, the measurement failed to obtain the correct optimal DRAG parameter for qubit q06, which should probably be closer to 0.1 instead of -0.057. A different or wider range of γ should have been swept instead.

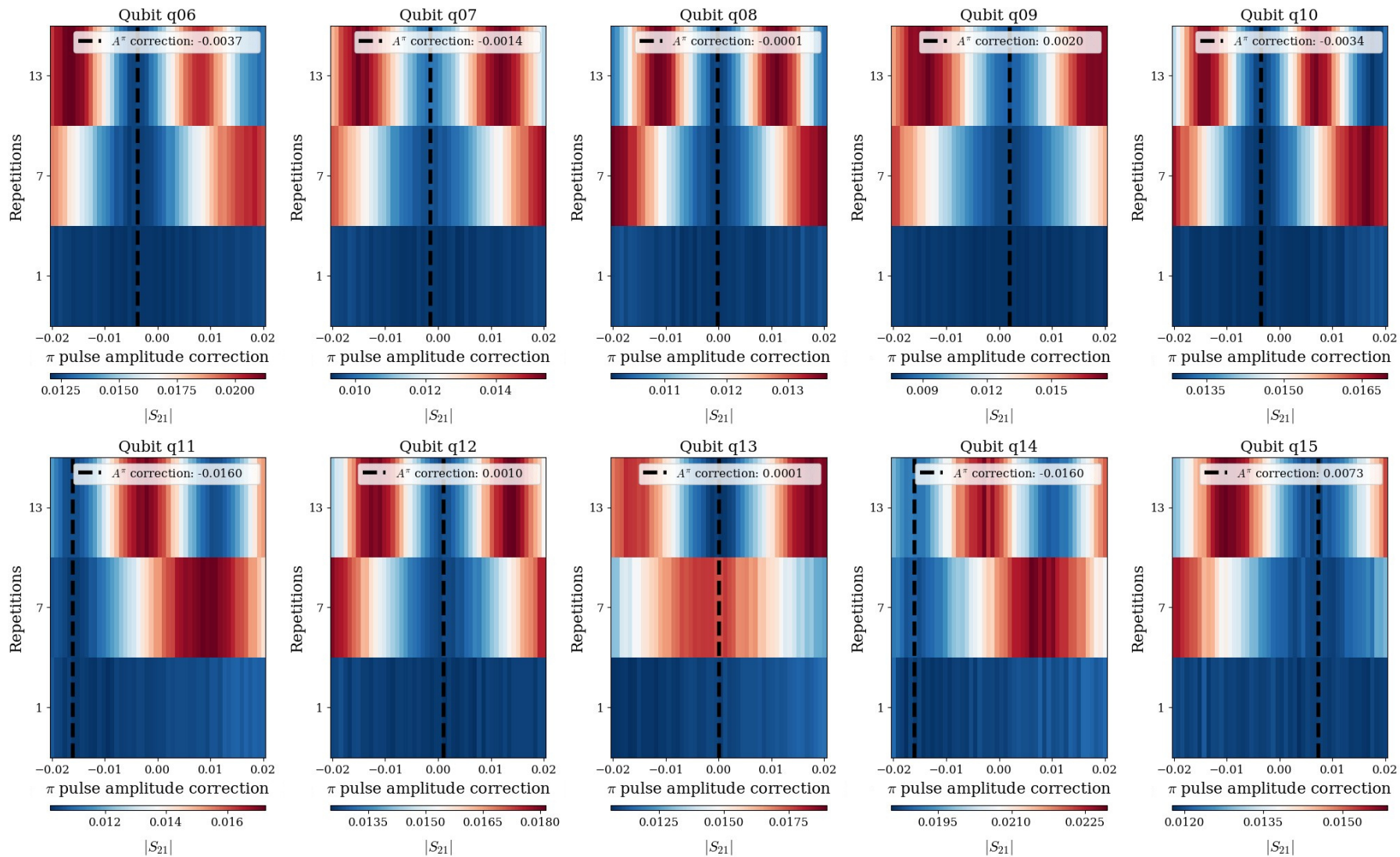


Figure A.11: n -Rabi oscillations for the $|0\rangle - |1\rangle$ transition. The black dashed line corresponds to the optimal π pulse amplitude correction. In this example, the correct value for qubit q13 lies outside the scanned amplitude range. A wider sweep towards more negative amplitudes should have been used instead for a better calibration.

A.3 Qubit characterisation nodes

A.3.1 T_1 measurement

This experiment is used to measure the relaxation time T_1 of the qubits. The qubit is first excited to state $|1\rangle$ and then allowed to idle for a variable delay time τ before being measured. During this delay, the qubit can spontaneously relax to its ground state $|0\rangle$. As the delay time increases, the probability of finding the qubit in the excited state decreases. Averaging the measurement outcomes over many repetitions yields an exponentially decaying excited-state population as a function of τ , from which the relaxation time T_1 can be extracted (Fig. A.12).

A.3.2 T_2 measurement

This measurement characterises the dephasing time T_2 of the qubits using the Ramsey sequence described in Section A.2.3. The qubit is prepared in a superposition state by the first $\pi/2$ pulse, after which it accumulates a relative phase during a variable delay time τ . Due to its interactions with the environment and fluctuations in the qubit frequency, phase coherence is gradually lost over time. As a result, the Ramsey fringes exhibit an exponentially decaying envelope and eventually converge to a population of 0.5, corresponding to a maximally mixed state. By fitting the decay envelope of the Ramsey oscillations, the dephasing time T_2 can be extracted.

A.3.3 T_2 Echo measurement

The T_2 Echo measurement is another technique used to characterise qubit dephasing while suppressing the effects of low-frequency noise and inhomogeneous broadening. The sequence consists of a $\pi/2$ pulse, followed by a delay time $\tau/2$, a refocusing π pulse, a second delay of $\tau/2$, and a final $\pi/2$ pulse before qubit measurement. The refocusing π pulse effectively reverses the phase accumulation caused by slowly varying frequency fluctuations, thereby mitigating dephasing from low-frequency noise sources. As the total delay time τ increases, the measured signal decays due to irreversible dephasing processes that cannot be refocused by the echo pulse. Fitting the decay envelope yields the spin-echo coherence time T_2^{Echo} (Fig. A.13). Since the low-frequency noise contribution is removed, T_2^{Echo} is generally longer than the Ramsey dephasing time T_2 .

Performing T_2 Echo measurements simultaneously on multiple qubits can reveal quantum crosstalk (see Section 1.5.2), as the refocusing π pulse in the sequence can only remove dephasing arising from local noise on each individual qubit. The phase accumulation caused by residual ZZ interaction between the qubits, however, cannot be fully refocused. As a result, qubits affected by quantum crosstalk have shorter T_2^{Echo} times when measured simultaneously than when measured individually, and their decay curves often develop non-exponential distortions, as can be seen for qubits q06 and q11 in Fig. A.13.

A.3.4 RB

RB is a widely used hardware characterisation protocol that provides an estimate of the average gate fidelity in a scalable manner while being insensitive to state preparation and measurement (SPAM) errors [58, 79]. For single-qubit gates, the protocol begins by initialising the qubit in its ground state. A sequence of m random Clifford gates¹ is then applied, collectively implementing a unitary operation C_m on the qubit. This is followed by a final gate, $C_m^\dagger$, to invert the net effect of the preceding gates, before the qubit is measured in its computational basis. The protocol is repeated for increasing numbers of Clifford gates m and for multiple different random sequences [79]. In the absence of errors, the qubit would always return to its initial state after $C_m^\dagger$. In practice, however, gate errors accumulate throughout the sequence, causing the probability of measuring the qubit in the ground state, P_0 , to decay exponentially with increasing sequence length (Figs. 4.4, 4.5, B.3, and B.4):

$$P_0 = A_0 p_0^m + B_0, \quad (\text{A.5})$$

where the constants A_0 and B_0 capture SPAM errors, while the decay parameter p_0 directly relates to the average error per Clifford gate ε_{avg} through

$$\varepsilon_{\text{avg}} = \frac{d-1}{d}(1-p_0), \quad (\text{A.6})$$

where d is the dimension of the Hilbert space [77]. For a single qubit, $d = 2$. The corresponding average Clifford gate fidelity is then given by

$$F_{\text{avg}} = 1 - \varepsilon_{\text{avg}} = \frac{1+p_0}{2}. \quad (\text{A.7})$$

This standard RB protocol can be extended to leakage RB to quantify leakage out of, and seepage back into, the computational subspace [64]. As for the ground state population P_0 , the computational-subspace population P_C can also be fitted to an exponentially decaying model:

$$P_C = A_1 p_1^m + B_1. \quad (\text{A.8})$$

The leakage rate L and seepage rate S are then extracted from the fitted parameters as:

$$L = (1 - B_1)(1 - p_1), \quad (\text{A.9})$$

$$S = B_1(1 - p_1). \quad (\text{A.10})$$

¹The Clifford group is a finite set of quantum gates that map Pauli operators onto other Pauli operators. For a single qubit, it consists of 24 gates generated by combinations of π and $\pi/2$ rotations about the X, Y, and Z axes of the Bloch sphere.

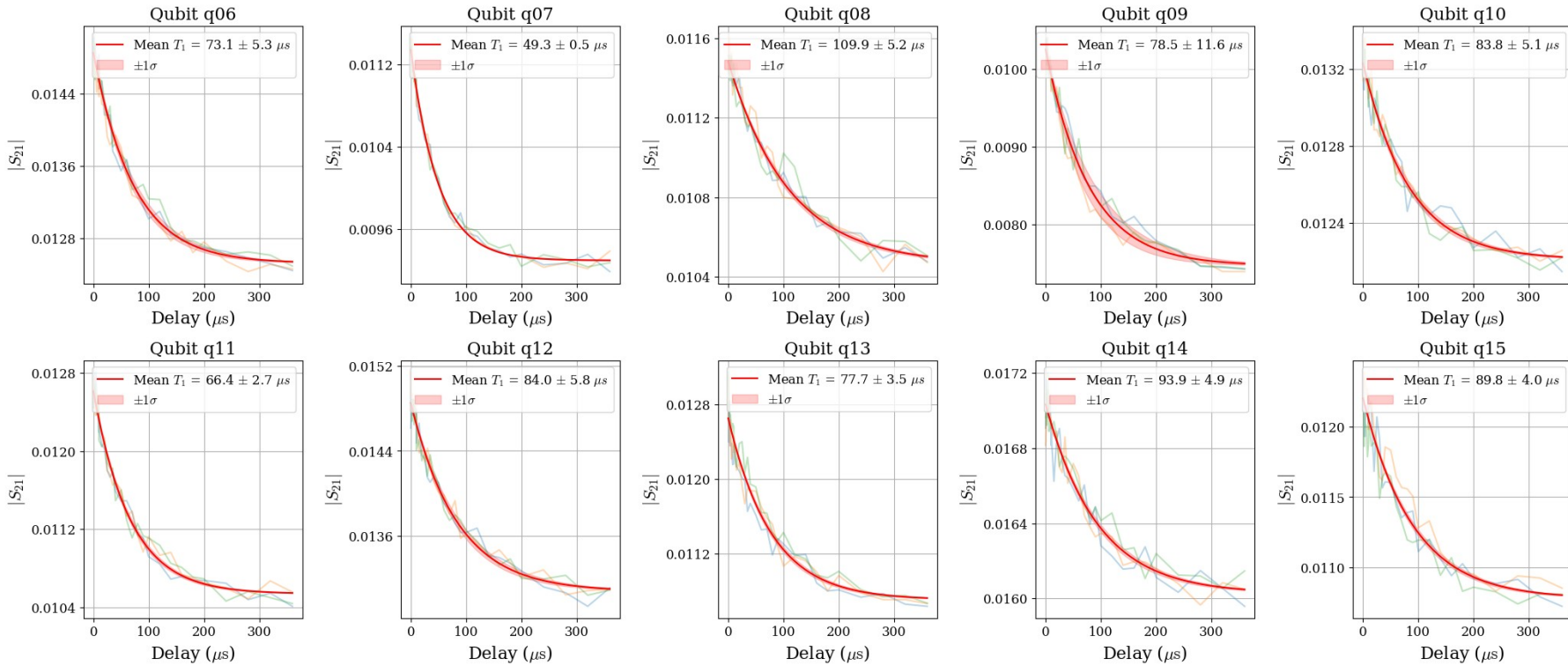


Figure A.12: T_1 measurement. The red curve is the average of three measurement repetitions (blue, green, and orange curves).

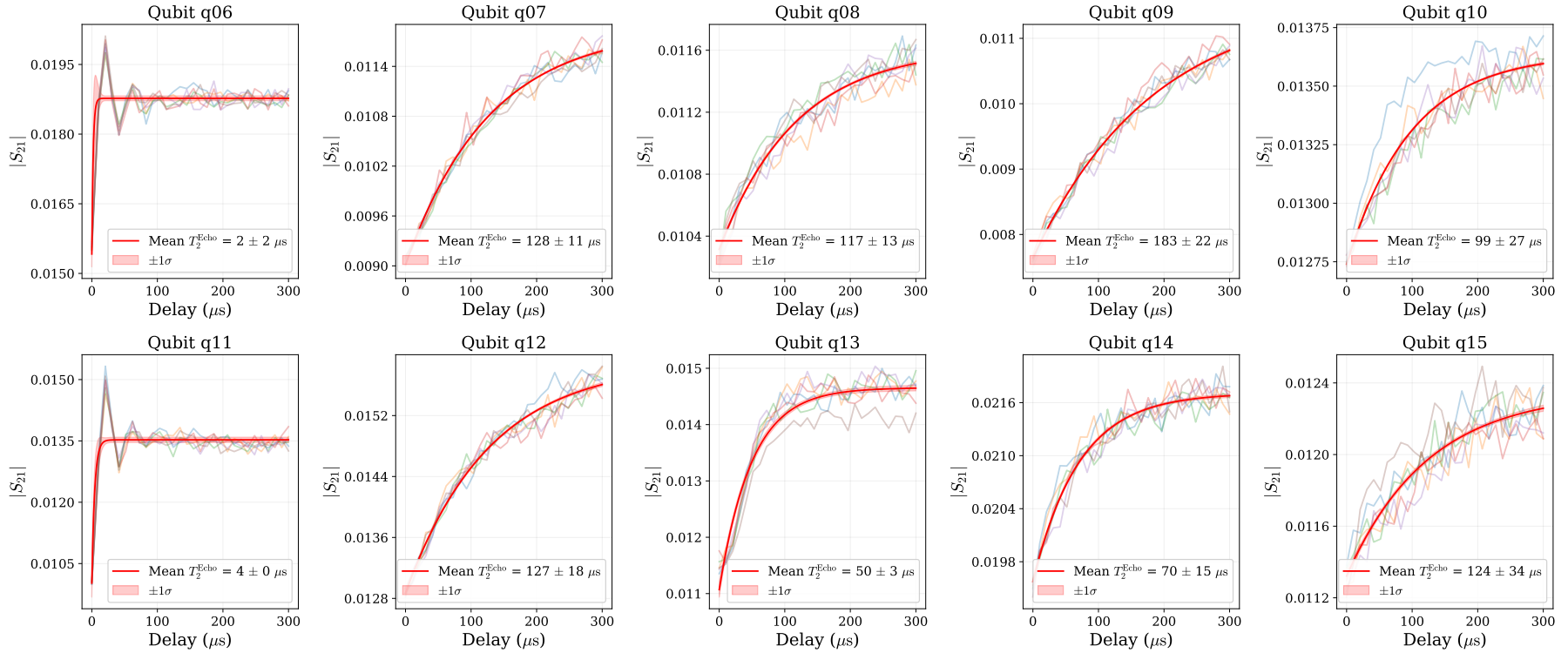


Figure A.13: T_2 Echo measurement. The red curve is the average of six measurement repetitions (the other curves). The unusually short T_2^{Echo} times and distorted decay curves observed for qubits q06 and q11 imply the presence of quantum crosstalk between the two qubits.

A.4 Crosstalk measurement nodes

These are the experiments developed and described in this thesis for the measurement of XY crosstalk amplitude and phase (see Chapter 4).

A.5 Example of a single-qubit gate calibration chain

This section presents an example single-qubit bring-up protocol applied to the 25-qubit flip-chip device. The exact parameter values are omitted as they may vary between qubits and between measurement sessions.

1. **Resonator spectroscopy 0:** Estimate the resonator frequency when the qubit is prepared in state $|0\rangle$.
2. **Punch-out 0:** Optimise the readout amplitude for state $|0\rangle$.
3. **Fine resonator spectroscopy 0:** Refine the resonator frequency using the optimal readout amplitude obtained from step 2.
4. **Two-dimensional qubit spectroscopy 01:** Estimate the first transition frequency of the qubit and the optimal qubit drive amplitude. After each cool-down, an additional preliminary one-dimensional frequency sweep over a wider range may be performed before this step to locate the transition frequency.
5. **Fine two-dimensional qubit spectroscopy 01:** Refine the transition frequency and drive amplitude from step 4 using narrower frequency and amplitude ranges.
6. **Rabi oscillations 01:** Determine the optimal π pulse amplitude for the $|0\rangle$ - $|1\rangle$ transition.
7. **Ramsey correction 01:** Fine-tune the qubit's first transition frequency found in step 5.
8. **DRAG parameter optimisation 01:** Determine the optimal DRAG pulse parameter for the $|0\rangle$ - $|1\rangle$ transition to minimise leakage to higher excited states.
9. **n-Rabi oscillations 01:** Refine the π pulse amplitude extracted from step 6.
10. **Resonator spectroscopy and punch-out 1:** Repeat steps 1 - 3 but with the qubit prepared in state $|1\rangle$.
11. **Qubit calibration 12:** Repeat steps 4 - 9 but for the $|1\rangle$ - $|2\rangle$ transition.
12. **Resonator spectroscopy and punch-out 2:** Repeat steps 1 - 3 but with the qubit prepared in state $|2\rangle$.
13. **Three-state readout frequency optimisation:** Optimise the resonator frequency for three-state SSRO.
14. **Three-state readout amplitude optimisation:** Optimise the readout amplitude for three-state SSRO.
15. **Single-qubit gate SSRO RB:** Extract the average single-qubit gate fidelity, leakage rate, and seepage rate.

Appendix B

Supplementary Materials

This appendix includes additional experimental results. Specifically, Figs. B.1 and B.2 are examples of $f_{01} - f_{01}$ crosstalk amplitude and phase results, respectively, obtained for qubits q11 and q15. Figs. B.3 and B.4 are detailed RB plots obtained with and without XY crosstalk compensation, respectively, for qubit q07, whose gate-error and leakage-rate matrices are shown in Section 4.3. Figs. B.5, B.6, and B.7 are average single-qubit gate error matrices for qubit pairs q11 - q13, q13 - q15, and q11 - q15, respectively.

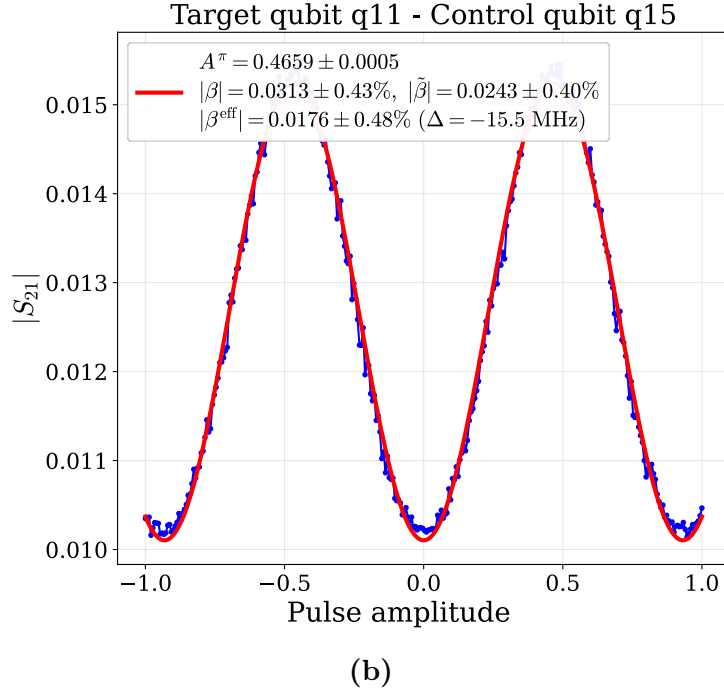
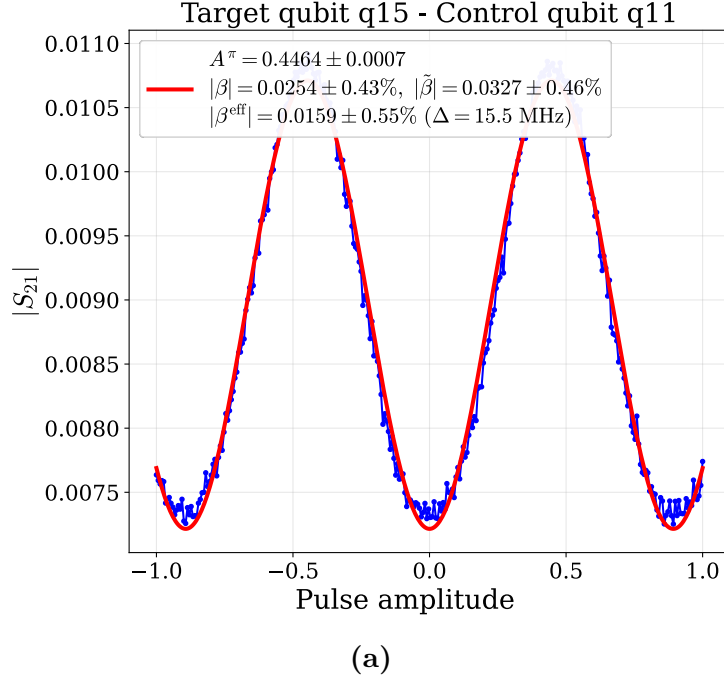
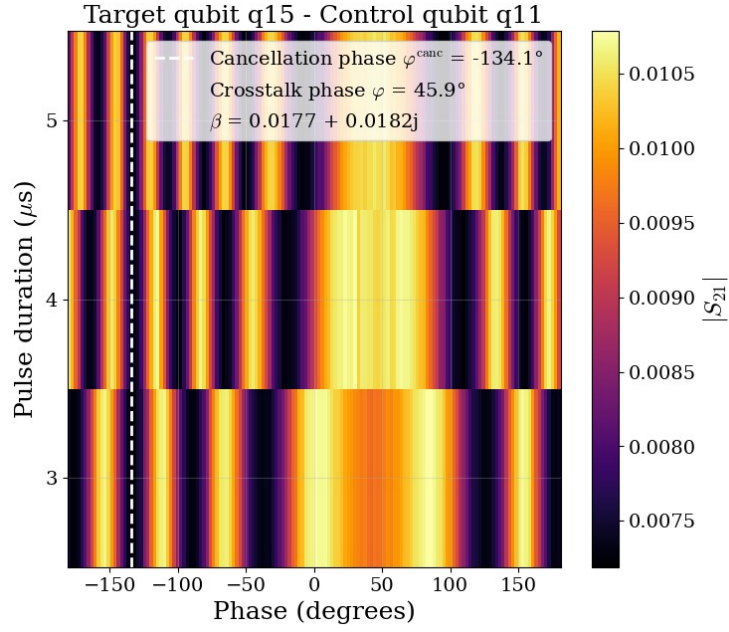
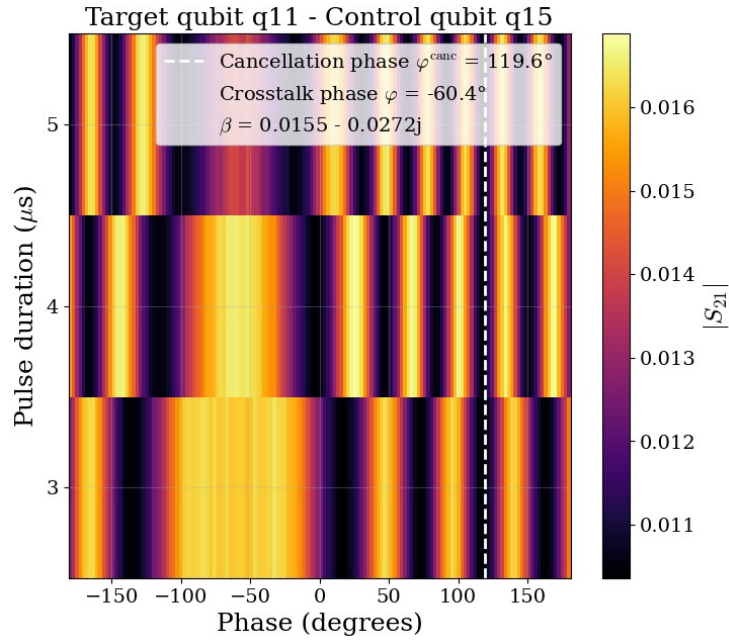


Figure B.1: Representative $f_{01} - f_{01}$ crosstalk amplitude results for qubits q11 and q15. (a) Qubit q11 acts as the control qubit and q15 as the target qubit. (b) Qubit q15 acts as the control qubit and q11 as the target qubit. Unlike $f_{01} - f_{12}$ crosstalk, $f_{01} - f_{01}$ crosstalk is symmetric, i.e., the effective crosstalk coefficients $|\beta^{\text{eff}}|$ are of similar amplitude in both directions.



(a)



(b)

Figure B.2: Representative $f_{01} - f_{01}$ crosstalk phase results for qubits q11 and q15. (a) Qubit q11 acts as the control qubit and q15 as the target qubit. (b) Qubit q15 acts as the control qubit and q11 as the target qubit. The white dashed line indicates the cancellation phase, at which the leaked tone and the compensation tone interfere destructively for all pulse durations. The crosstalk phase is 180° out of phase with the cancellation phase. The complex crosstalk coefficient β shown in each panel is calculated from the extracted crosstalk phase together with the amplitude $|\beta|$ obtained from the amplitude measurement (Fig. B.1).

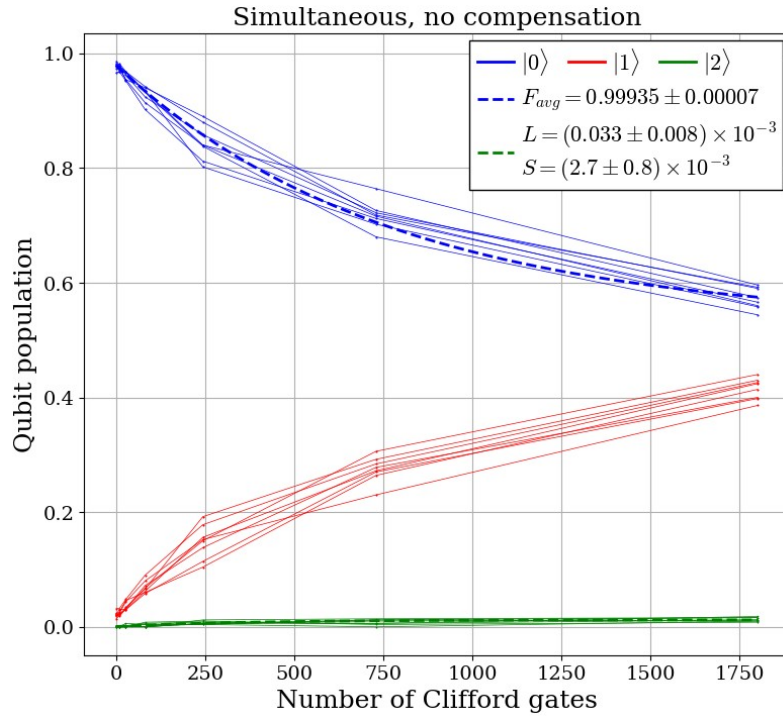
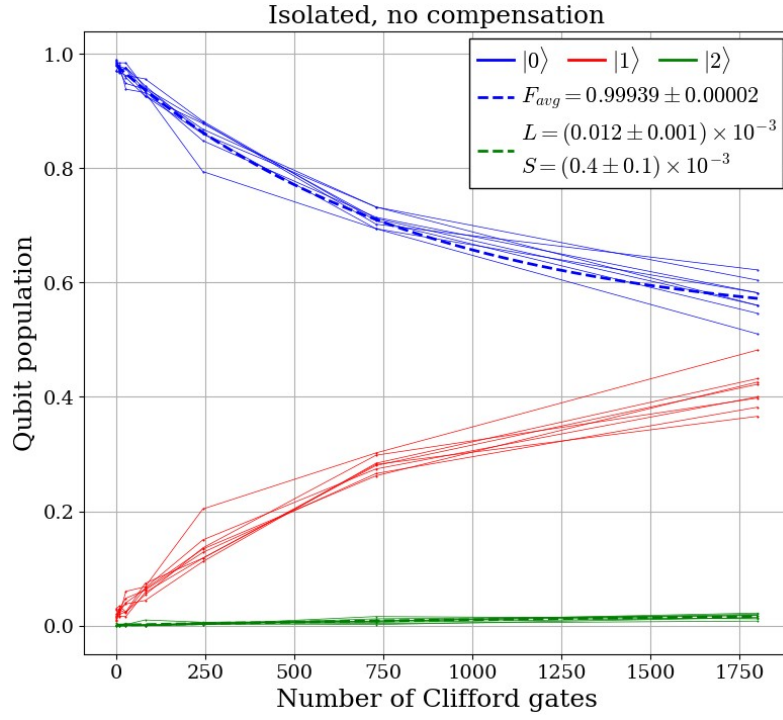
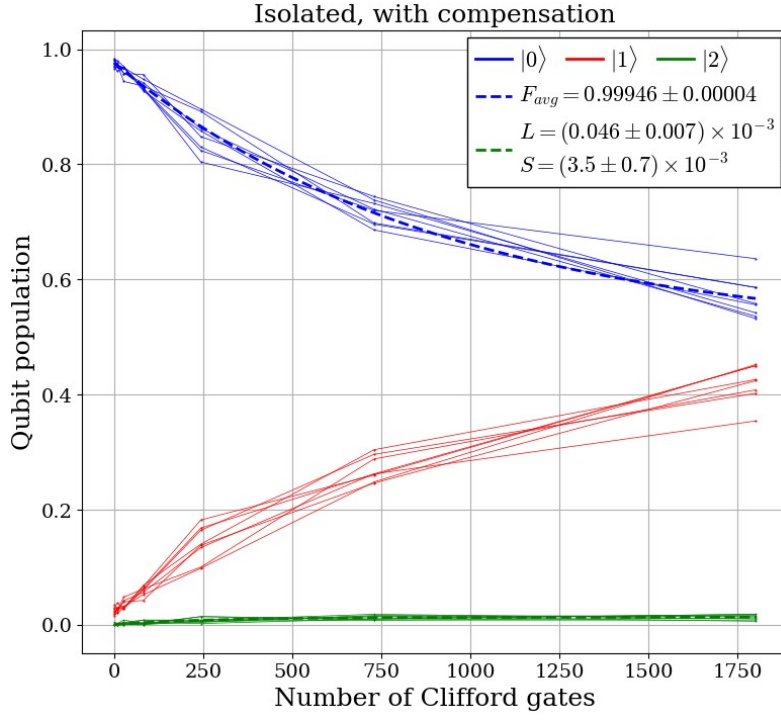
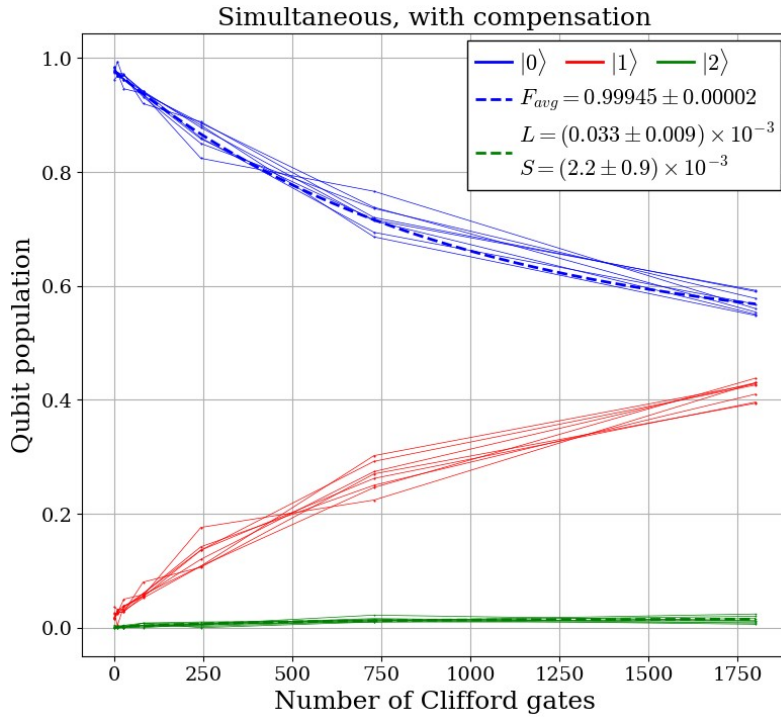


Figure B.3: RB results for qubit q07 without crosstalk compensation. (a) Isolated RB. (b) Simultaneous RB performed with qubit q09.



(a)



(b)

Figure B.4: RB results for qubit q07 with crosstalk compensation. (a) Isolated RB. (b) Simultaneous RB performed with qubit q09.

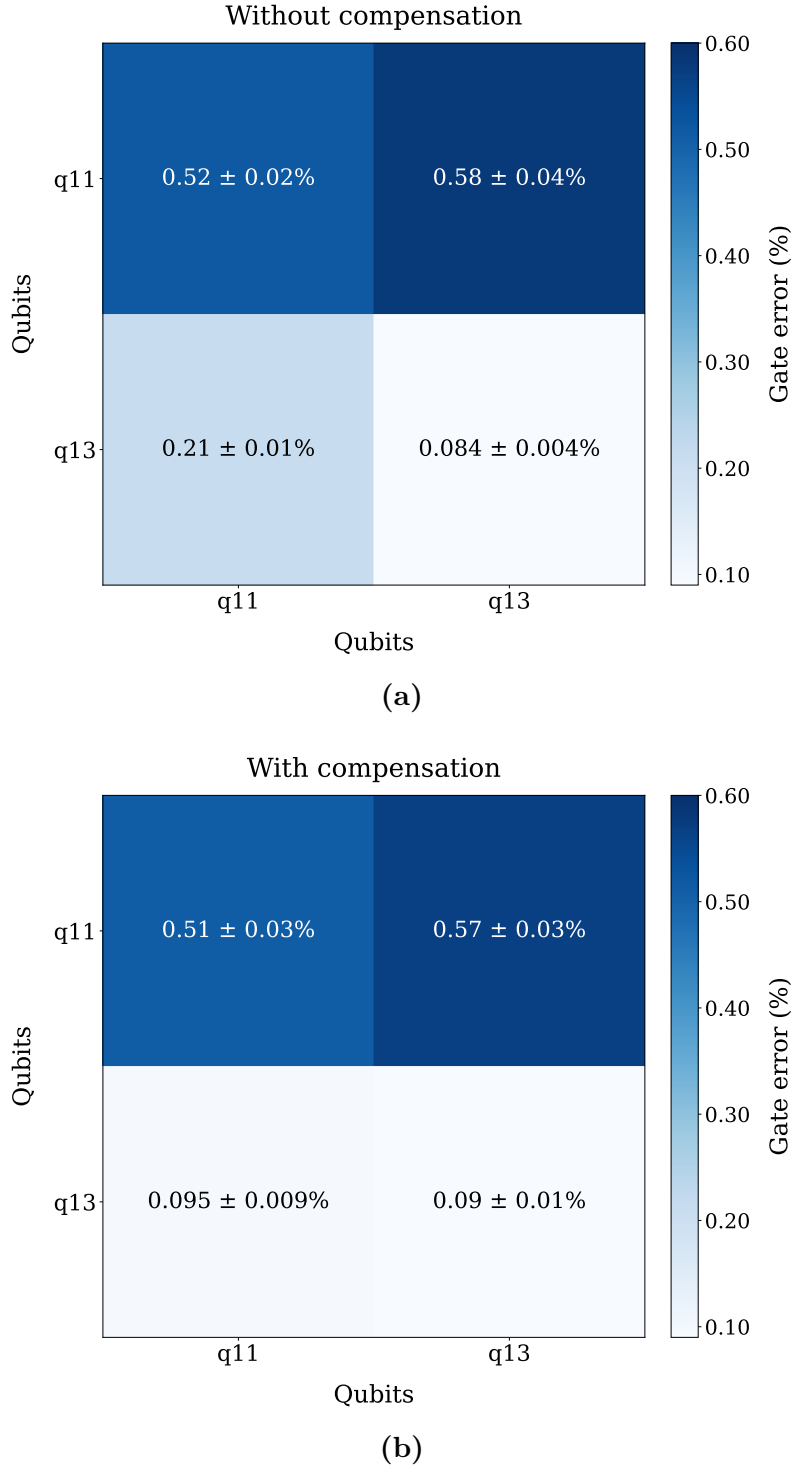
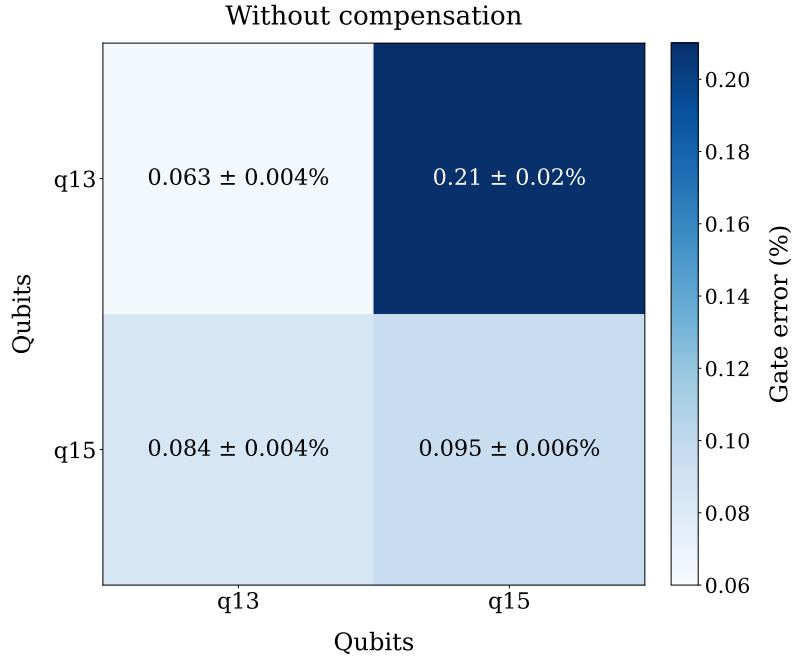
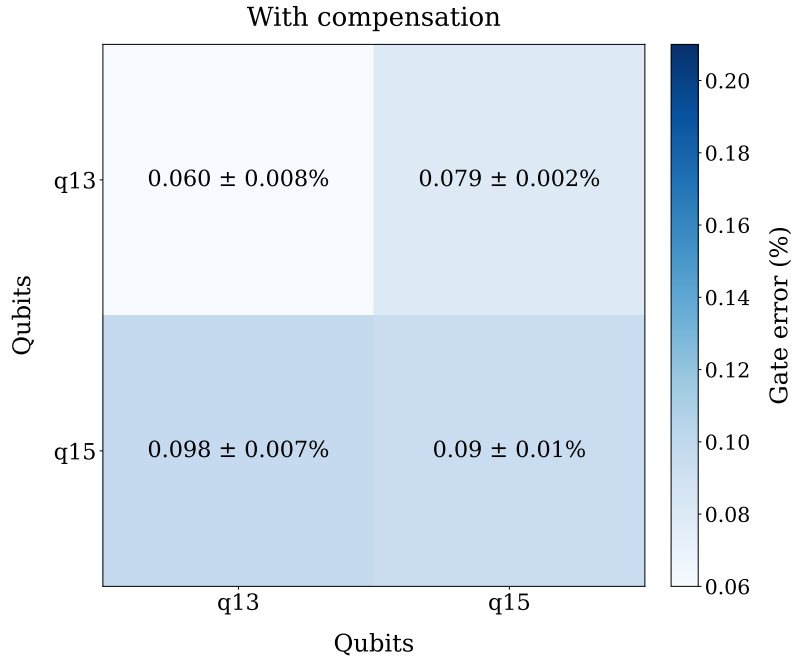


Figure B.5: Average single-qubit gate error matrix for qubits q11 and q13. (a) Without crosstalk compensation, qubit q13's gate error increases by approximately 2.5 times during simultaneous operation. (b) With crosstalk compensation, qubit q13's gate errors obtained from isolated and simultaneous RB agree within experimental uncertainties.



(a)



(b)

Figure B.6: Average single-qubit gate error matrix for qubits q13 and q15. (a) Without crosstalk compensation, qubit q13's gate error increases by approximately 3.3 times during simultaneous operation. (b) With crosstalk compensation, qubit q13's gate error during simultaneous RB is reduced to only 1.3 times higher than the isolated value.

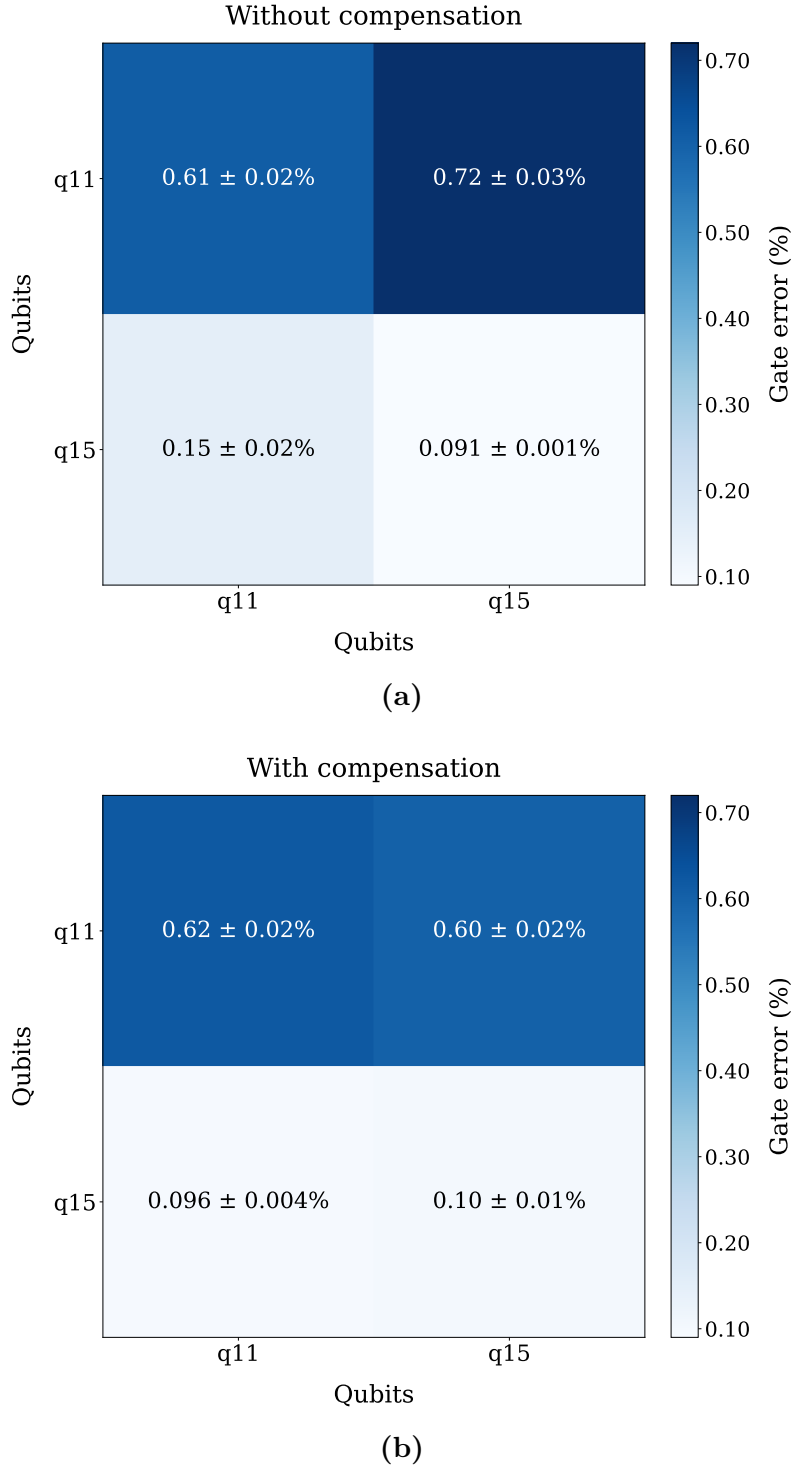


Figure B.7: Average single-qubit gate error matrix for qubits q11 and q15. (a) Without crosstalk compensation, both qubit q11 and qubit q15's gate errors increase significantly during simultaneous operations. (b) With crosstalk compensation, the gate errors obtained for each qubit from isolated and simultaneous RB agree within experimental uncertainties.

Bibliography

- [1] Richard P Feynman. “Simulating physics with computers”. In: *Feynman and computation*. cRc Press, 2018, pp. 133–153.
- [2] Paul Benioff. “The computer as a physical system: A microscopic quantum mechanical Hamiltonian model of computers as represented by Turing machines”. In: *Journal of statistical physics* 22.5 (1980), pp. 563–591.
- [3] Paul Benioff. “Quantum mechanical models of Turing machines that dissipate no energy”. In: *Physical Review Letters* 48.23 (1982), p. 1581.
- [4] David Deutsch. “Quantum theory, the Church–Turing principle and the universal quantum computer”. In: *Proceedings of the Royal Society of London. A. Mathematical and Physical Sciences* 400.1818 (1985), pp. 97–117.
- [5] Daniel R Simon. “On the power of quantum computation”. In: *SIAM journal on computing* 26.5 (1997), pp. 1474–1483.
- [6] Peter W Shor. “Polynomial-time algorithms for prime factorization and discrete logarithms on a quantum computer”. In: *SIAM review* 41.2 (1999), pp. 303–332.
- [7] Lov K Grover. “A fast quantum mechanical algorithm for database search”. In: *Proceedings of the twenty-eighth annual ACM symposium on Theory of computing*. 1996, pp. 212–219.
- [8] Muhammad AbuGhanem. “IBM quantum computers: evolution, performance, and future directions: M. AbuGhanem”. In: *The Journal of Supercomputing* 81.5 (2025), p. 687.
- [9] Muhammad AbuGhanem. “Superconducting quantum computers: who is leading the future?” In: *EPJ Quantum Technology* 12.1 (2025), p. 102.
- [10] David P DiVincenzo. “The physical implementation of quantum computation”. In: *Fortschritte der Physik: Progress of Physics* 48.9-11 (2000), pp. 771–783.
- [11] Colin D Bruzewicz et al. “Trapped-ion quantum computing: Progress and challenges”. In: *Applied physics reviews* 6.2 (2019).
- [12] John Preskill. “Quantum computing 40 years later”. In: *Feynman lectures on computation*. CRC Press, 2023, pp. 193–244.
- [13] David S Weiss and Mark Saffman. “Quantum computing with neutral atoms”. In: *Physics Today* 70.7 (2017), pp. 44–50.

- [14] Guido Burkard et al. “Semiconductor spin qubits”. In: *Reviews of Modern Physics* 95.2 (2023), p. 025003.
- [15] Hui Wang et al. “Scalable photonic quantum technologies”. In: *Nature Materials* 24.12 (2025), pp. 1883–1897.
- [16] Yao-Yao Jiang et al. “Advancements in superconducting quantum computing”. In: *National Science Review* 12.8 (2025), nwaf246.
- [17] David Aasen et al. “Roadmap to fault tolerant quantum computation using topological qubit arrays”. In: *arXiv preprint arXiv:2502.12252* (2025).
- [18] Rami Barends et al. “Superconducting quantum circuits at the surface code threshold for fault tolerance”. In: *Nature* 508.7497 (2014), pp. 500–503.
- [19] Austin G Fowler et al. “Surface codes: Towards practical large-scale quantum computation”. In: *Physical Review A—Atomic, Molecular, and Optical Physics* 86.3 (2012), p. 032324.
- [20] “Quantum error correction below the surface code threshold”. In: *Nature* 638.8052 (2025), pp. 920–926.
- [21] John M Martinis, Michel H Devoret, and John Clarke. “Energy-level quantization in the zero-voltage state of a current-biased Josephson junction”. In: *Physical review letters* 55.15 (1985), p. 1543.
- [22] John M Martinis, Michel H Devoret, and John Clarke. “Experimental tests for the quantum behavior of a macroscopic degree of freedom: The phase difference across a Josephson junction”. In: *Physical Review B* 35.10 (1987), p. 4682.
- [23] Yasunobu Nakamura, Yu A Pashkin, and JS Tsai. “Coherent control of macroscopic quantum states in a single-Cooper-pair box”. In: *nature* 398.6730 (1999), pp. 786–788.
- [24] Denis Vion et al. “Manipulating the quantum state of an electrical circuit”. In: *Science* 296.5569 (2002), pp. 886–889.
- [25] Irinel Chiorescu et al. “Coherent quantum dynamics of a superconducting flux qubit”. In: *Science* 299.5614 (2003), pp. 1869–1871.
- [26] John M Martinis et al. “Rabi oscillations in a large Josephson-junction qubit”. In: *Physical review letters* 89.11 (2002), p. 117901.
- [27] Alexandre Blais et al. “Circuit quantum electrodynamics”. In: *Reviews of Modern Physics* 93.2 (2021), p. 025005.
- [28] Jens Koch et al. “Charge-insensitive qubit design derived from the Cooper pair box”. In: *Physical Review A—Atomic, Molecular, and Optical Physics* 76.4 (2007), p. 042319.
- [29] Rami Barends et al. “Coherent Josephson qubit suitable for scalable quantum integrated circuits”. In: *Physical review letters* 111.8 (2013), p. 080502.
- [30] Vladimir E Manucharyan et al. “Fluxonium: Single cooper-pair circuit free of charge offsets”. In: *Science* 326.5949 (2009), pp. 113–116.

-
- [31] Atharv Joshi, Kyungjoo Noh, and Yvonne Y Gao. “Quantum information processing with bosonic qubits in circuit QED”. In: *Quantum Science & Technology* 6.3 (2021), p. 033001.
 - [32] Yulin Chi et al. “A programmable qudit-based quantum processor”. In: *Nature communications* 13.1 (2022), p. 1166.
 - [33] Theodore Van Duzer and Charles W Turner. “Principles of superconducting devices and circuits”. In: *Edward Arnold, London United Kingdom* (1981).
 - [34] F Benatti et al. “A quantum fluctuation description of charge qubits”. In: *New Journal of Physics* 26.1 (2024), p. 013057.
 - [35] Philip Krantz et al. “A quantum engineer’s guide to superconducting qubits”. In: *Applied physics reviews* 6.2 (2019).
 - [36] Easwar Magesan and Jay M Gambetta. “Effective Hamiltonian models of the cross-resonance gate”. In: *Physical Review A* 101.5 (2020), p. 052308.
 - [37] J Bergli, Yu M Galperin, and BL Altshuler. “Decoherence in qubits due to low-frequency noise”. In: *New Journal of Physics* 11.2 (2009), p. 025002.
 - [38] Gianluigi Catelani et al. “Decoherence of superconducting qubits caused by quasiparticle tunneling”. In: *Physical Review B—Condensed Matter and Materials Physics* 86.18 (2012), p. 184514.
 - [39] Chen Wang et al. “Measurement and control of quasiparticle dynamics in a superconducting qubit”. In: *Nature communications* 5.1 (2014), p. 5836.
 - [40] Diego Ristè et al. “Millisecond charge-parity fluctuations and induced decoherence in a superconducting transmon qubit”. In: *Nature communications* 4.1 (2013), p. 1913.
 - [41] John M Martinis et al. “Decoherence in Josephson qubits from dielectric loss”. In: *Physical review letters* 95.21 (2005), p. 210503.
 - [42] Robert McDermott. “Materials origins of decoherence in superconducting qubits”. In: *IEEE Transactions on Applied Superconductivity* 19.1 (2009), pp. 2–13.
 - [43] Antti P Vepsäläinen et al. “Impact of ionizing radiation on superconducting qubit coherence”. In: *Nature* 584.7822 (2020), pp. 551–556.
 - [44] Yoshiki Sunada et al. “Fast readout and reset of a superconducting qubit coupled to a resonator with an intrinsic Purcell filter”. In: *Physical Review Applied* 17.4 (2022), p. 044016.
 - [45] Eyob A Sete, John M Martinis, and Alexander N Korotkov. “Quantum theory of a bandpass Purcell filter for qubit readout”. In: *Physical Review A* 92.1 (2015), p. 012325.
 - [46] David A Rower et al. “Evolution of $1/f$ flux noise in superconducting qubits with weak magnetic fields”. In: *Physical Review Letters* 130.22 (2023), p. 220602.
 - [47] Amr Osman. *Scaling Superconducting Quantum Processors Coherence, Frequency Targeting and Crosstalk*. Chalmers Tekniska Högskola (Sweden), 2024.
 - [48] Sandoko Kosen et al. “Building blocks of a flip-chip integrated superconducting quantum processor”. In: *Quantum Science & Technology* 7.3 (2022), p. 035018.

- [49] A Dunsworth et al. “A method for building low loss multi-layer wiring for superconducting microwave devices”. In: *Applied Physics Letters* 112.6 (2018).
- [50] Sandoko Kosen et al. “Signal crosstalk in a flip-chip quantum processor”. In: *PRX quantum* 5.3 (2024), p. 030350.
- [51] D Rosenberg et al. “3D integrated superconducting qubits”. In: *npj quantum information* 3.1 (2017), p. 42.
- [52] CR Conner et al. “Superconducting qubits in a flip-chip architecture”. In: *Applied Physics Letters* 118.23 (2021).
- [53] Xuegang Li et al. “Vacuum-gap transmon qubits realized using flip-chip technology”. In: *Applied Physics Letters* 119.18 (2021).
- [54] Wuerkaixi Nuerbolati et al. “Canceling microwave crosstalk with fixed-frequency qubits”. In: *Applied Physics Letters* 120.17 (2022).
- [55] Jan Balewski et al. “First-principle crosstalk dynamics and Hamiltonian learning via Rabi experiments”. In: *Proceedings of the 23rd Annual International Conference on Mobile Systems, Applications and Services*. 2025, pp. 765–770.
- [56] Ruixia Wang et al. “Control and mitigation of microwave crosstalk effect with superconducting qubits”. In: *Applied Physics Letters* 121.15 (2022).
- [57] Andreas Ketterer and Thomas Wellens. “Characterizing crosstalk of superconducting transmon processors”. In: *Physical Review Applied* 20.3 (2023), p. 034065.
- [58] Jay M Gambetta et al. “Characterization of addressability by simultaneous randomized benchmarking”. In: *Physical review letters* 109.24 (2012), p. 240504.
- [59] Seungchan Seo and Joonwoo Bae. “Measurement crosstalk errors in cloud-based quantum computing”. In: *IEEE Internet Computing* 26.1 (2021), pp. 26–33.
- [60] Fabio Altomare et al. “Measurement crosstalk between two phase qubits coupled by a coplanar waveguide”. In: *Physical Review B—Condensed Matter and Materials Physics* 82.9 (2010), p. 094510.
- [61] Peng Zhao et al. “Spurious microwave crosstalk in floating superconducting circuits”. In: *arXiv preprint arXiv:2206.03710* (2022).
- [62] James Wenner et al. “Wirebond crosstalk and cavity modes in large chip mounts for superconducting qubits”. In: *Superconductor Science and Technology* 24.6 (2011), p. 065001.
- [63] Yongxin Song et al. “Microwave Crosstalk in Planar Superconducting Quantum Devices”. In: *arXiv preprint arXiv:2606.02440* (2026).
- [64] Christopher J Wood and Jay M Gambetta. “Quantification and characterization of leakage errors”. In: *Physical Review A* 97.3 (2018), p. 032306.
- [65] Austin G Fowler. “Coping with qubit leakage in topological codes”. In: *Physical Review A—Atomic, Molecular, and Optical Physics* 88.4 (2013), p. 042308.
- [66] Felix Motzoi et al. “Simple pulses for elimination of leakage in weakly nonlinear qubits”. In: *Physical review letters* 103.11 (2009), p. 110501.

- [67] Xiao-Yan Yang et al. “Fast, universal scheme for calibrating microwave crosstalk in superconducting circuits”. In: *Applied Physics Letters* 125.4 (2024).
- [68] Haisheng Yan et al. “Calibration and cancellation of microwave crosstalk in superconducting circuits”. In: *Chinese Physics B* 32.9 (2023), p. 094203.
- [69] Ryo Matsuda et al. “Selective excitation of superconducting qubits with a shared control line through pulse shaping”. In: *Physical Review Research* 8.1 (2026), p. L012003.
- [70] Jaap J Wesdorp et al. “Mitigating crosstalk errors for simultaneous single-qubit gates on a superconducting quantum processor”. In: *arXiv preprint arXiv:2603.11018* (2026).
- [71] ZH Yang et al. “Mitigation of microwave crosstalk with parameterized single-qubit gate in superconducting quantum circuits”. In: *Applied Physics Letters* 124.21 (2024).
- [72] Jay M Gambetta et al. “Analytic control methods for high-fidelity unitary operations in a weakly nonlinear oscillator”. In: *Physical Review A—Atomic, Molecular, and Optical Physics* 83.1 (2011), p. 012308.
- [73] Amr Osman et al. “Mitigation of frequency collisions in superconducting quantum processors”. In: *Physical Review Research* 5.4 (2023), p. 043001.
- [74] Christian Križan et al. “Electrical post-fabrication tuning of aluminum Josephson junctions at room temperature”. In: *Materials for Quantum Technology* (2026).
- [75] Qblox. *Cluster*. 2025. URL: <https://qblox.com/product/cluster> (visited on 06/15/2026).
- [76] *Tergite Autocalibration*. 2026. URL: <https://github.com/tergite/tergite-autocalibration> (visited on 08/13/2026).
- [77] Joel Sandås. “Toward an automatic parallelized two-qubit-gate calibration”. In: (2025).
- [78] Zijun Chen. “Metrology of quantum control and measurement in superconducting qubits”. PhD thesis. University of California, Santa Barbara, 2018.
- [79] Emanuel Knill et al. “Randomized benchmarking of quantum gates”. In: *Physical Review A—Atomic, Molecular, and Optical Physics* 77.1 (2008), p. 012307.

Code of Conduct & Transparency Statement on the Use of GenAI for KU Leuven Students (Academic Year 2025–2026)

Generative AI (GenAI) assistance tools can be used to generate various types of content, including text, images, code, video, music, or combinations thereof. Common examples of such tools include ChatGPT, Google Gemini, Microsoft Copilot, Midjourney, Claude.ai, Perplexity.ai, and DALL·E, among others.

This code of conduct is a tool that helps students to be transparent about the use of GenAI and fits within the university's principles on academic integrity.

Important guidelines and remarks

- **Sensitive or personal data:** Some GenAI tools protect your input and sensitive or personal data better than others. There is often no transparency on what the owners of the AI applications do with the data entered. Therefore, **do not enter sensitive or personal data in free GenAI tools.** More info about what sensitive and personal data are, is described in the three confidentiality levels of the KU Leuven data classification model (non-confidential, confidential, strictly confidential). For confidential data you should preferably use **M365 Copilot Chat with your KU Leuven account.** If you use another GenAI tool, you must be **absolutely certain** it does **not store or reuse the data you enter.** Try to avoid entering these data. Do so only if strictly necessary, and only to the extent required. For personal data, work with anonymous or pseudonymised data. Additionally, in case of strictly confidential data this should first be discussed with the teaching staff of the course or your thesis supervisor.
- **Copyrighted materials:** For lawfully obtained copyrighted material you should preferably use **M365 Copilot Chat with your KU Leuven account.** If you use another GenAI tool, you must be **absolutely certain** it does **not store or reuse the data you enter.**
- GenAI assistance may not be used for data or topics covered by a Non-Disclosure Agreement (NDA). Even in the absence of an NDA, certain information may still need to be treated as confidential—for example, due to regulatory requirements or the risk of significant harm to the university if

disclosed. In case of doubt, check with your teaching staff or supervisor.

- If your master's thesis is under embargo, you should first discuss with your supervisor whether the use of GenAI is permitted.
- Before using a GenAI tool, always consider whether its use is responsible, including from a sustainability perspective (e.g. when using GenAI as a search engine, language assistant, ...).
- **Take a scientific and critical attitude** when interacting with GenAI assistance and interpreting its output, that may not always be correct.
- As a student you are responsible for complying with Article 84 of the Regulations on Education and Examinations: your report or thesis should reflect your own knowledge, understanding and skills. Be aware that plagiarism rules also apply to (work that is the result of) the use of GenAI assistance tools.

Exam Regulations Article 84: *“Every conduct individual students display with which they (partially) inhibit or attempt to inhibit a correct judgement of their own knowledge, understanding and/or skills or those of other students, is considered an irregularity which may result in a suitable penalty. A special type of irregularity is plagiarism, i.e. copying the work (ideas, texts, structures, designs, images, plans, codes , ...) of others or prior personal work in an exact or slightly modified way without adequately acknowledging the sources. Every possession of prohibited resources during an examination (see article 65) is considered an irregularity.”*

- In order to maintain academic integrity and avoid plagiarism, **more information about being transparent on the use of GenAI assistance and about correctly citing and referencing GenAI** can be found on this website for students (Dutch/English).
- **Additional reading: KU Leuven guidelines on responsible use of Generative AI tools, and other information** (Dutch/English)

A few final words

If you are uncertain whether or not you should declare your use of GenAI tools, we suggest that you discuss this with your instructor or supervisor. It is always safer to declare GenAI use, even when it is not strictly required. However, declaring GenAI use does not entail that its use is allowed; the right column in the table below provides more detailed instructions in this regard (code of conduct).

Moreover, advanced AI tools are evolving rapidly, and their capabilities have expanded significantly in a short period of time. As a result, we do not yet have all the answers about their responsible use. Finally, it is important to follow-up on the most recent evolutions in AI technologies, to have an open mind but also to be a bit cautious, to communicate with instructors, teaching assistants, supervisors and peers, to be as transparent as we can, and to learn together as we move along.

Student name: Hong Quan Trinh

Student number: 1019315

Please indicate with "**X**" whether it relates to a course assignment or to the Bachelor's or Master's thesis:

☐ This form is related to a **course assignment**.

Course name: **Course number:**

This form is related to my ☐ **Bachelor's** or ☒ **Master's thesis**.

Title: XY Crosstalk in a 25-Qubit Flip-Chip Superconducting Quantum Processor

Supervisor: Dr. Michele Faucci Giannelli

Please indicate with "**X**":

☐ I **did not use** any GenAI assistance tool.

☒ I **did use** GenAI Assistance. In this case **specify which ones** (e.g. ChatGPT, M365 Copilot, ...): ChatGPT and Claude.

GenAI assistance used as/for:	Name of the GenAI tool(s) used If helpful, also describe in which way you were using GenAI related to what is specified as code of conduct.	Code of conduct: For each of the categories below, always take into account the important guidelines and remarks mentioned above (e.g. copyrighted data, sensitive or personal data,...).
As a language assistant for reviewing or improving texts I wrote myself	<i>ChatGPT</i>	This use is similar to using spelling and grammar check tools. In general, you do not have to refer to such kind of GenAI use in the text. However, be careful: When using GenAI tools on texts you did not write yourself to improve the text, you have to refer to the original source or author, otherwise you are committing plagiarism and thus an irregularity.
As a paraphrasing tool	<i>ChatGPT</i>	This use is allowed except when it is prohibited by the teacher or the program of study. You may paraphrase your own text or texts by an author other than yourself and take inspiration from what a GenAI tool or another tool suggests (unless it is not allowed). In general, you do not have to refer to such kind of GenAI use or other paraphrasing tools in the text. However, be careful: If it entails text by an author other than yourself, you are not allowed to include that paraphrased text without reference to the original source or author. Without such reference, you would be committing plagiarism and thus an irregularity.

As a search engine to get information on a topic or to search for existing research on the topic	<i>ChatGPT, Claude</i>	This use is similar to e.g. a Google search or checking Wikipedia. If you write your own text based on this information, you do not have to refer to the use of GenAI in the text. You only have to refer to the existing research and references you have checked and used (without such references you would be committing plagiarism and thus committing an irregularity). However, be careful: • Be aware that the output of the GenAI tool cannot be guaranteed as a 100% reliable source of information. The output may not be entirely correct and/or be limited due to the databases it uses. Moreover, knowledge evolves and may change over time; therefore, the database of the GenAI tool may not be up to date. Therefore, verify the information and do not just copy-paste it as you should understand and critically process everything you are writing.
For generating programming code	<i>Claude</i>	Use of GenAI for coding is allowed except when it is prohibited by the teacher or the program of study. If used for coding, correctly mention the use of GenAI assistance in accordance with the instructions on the page on being transparent about the use of GenAI.
For generating new (research) ideas	<i>ChatGPT, Claude</i>	This use of GenAI is allowed except when it is prohibited by the teacher or the program of study. Correctly mention the use of GenAI assistance in accordance with the instructions on the page on being transparent about the use of GenAI. Be careful: • Further verify in this case whether the idea is novel or not. It is likely that it is related to existing work. If so, that existing work should be correctly referenced in the text (without such reference you would be committing plagiarism and thus committing an irregularity).

