

Quantum error correction with superconducting qubits

Ferreira Marques, J.M.

DOI

[10.4233/uuid:24e8c212-1d5a-4ecc-a520-27030686607b](https://doi.org/10.4233/uuid:24e8c212-1d5a-4ecc-a520-27030686607b)

Publication date

2024

Document Version

Final published version

Citation (APA)

Ferreira Marques, J. M. (2024). *Quantum error correction with superconducting qubits*. [Dissertation (TU Delft), Delft University of Technology]. <https://doi.org/10.4233/uuid:24e8c212-1d5a-4ecc-a520-27030686607b>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

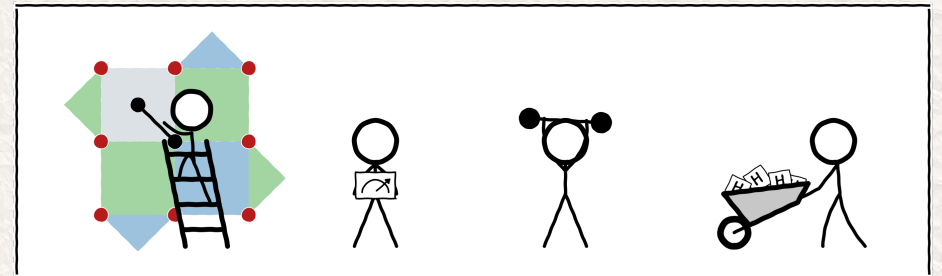
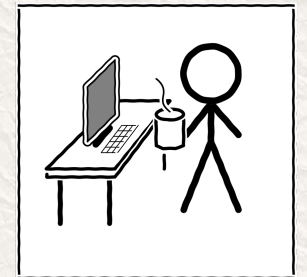
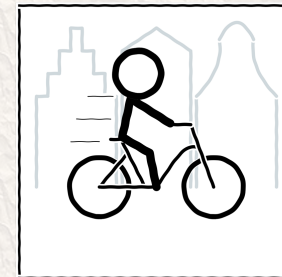
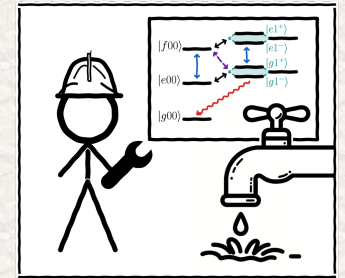
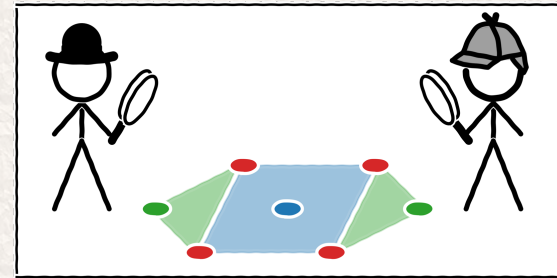
Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

QUANTUM ERROR CORRECTION WITH SUPERCONDUCTING QUBITS



QUANTUM ERROR CORRECTION WITH SUPERCONDUCTING QUBITS

QUANTUM ERROR CORRECTION WITH SUPERCONDUCTING QUBITS

Dissertation

for the purpose of obtaining the degree of doctor
at Delft University of Technology
by the authority of the Rector Magnificus Prof. Dr. Ir. T.H.J.J. van der Hagen,
Chair of the Board for Doctorates,
to be defended publicly on 27 September 2024, 10:00 o'clock.

by

Jorge MIGUEL FERREIRA MARQUES

Master of Science in Engineering Physics,
Insituto Superior Técnico, Portugal,
born in Porto, Portugal.

This dissertation has been approved by the promotor.

Composition of the doctoral committee:

Rector Magnificus,	Chairperson
Prof. dr. L. DiCarlo,	Delft University of Technology, promotor
Prof. dr. B. M. Terhal,	Delft University of Technology, promotor

Independent members:

Dr. C. K. Andersen	Delft University of Technology
Prof. dr. ir. L. P. Kouwenhoven	Delft University of Technology
Prof. dr. M. Müller	RWTH Aachen University, Germany
Prof. dr. V. E. Manucharyan	Ecole Polytechnique Federale de Lausanne, Switzerland
Prof. dr. ir. R. Hanson	Delft University of Technology, Reserve member



Printed by: Gildeprint, Enschede – www.gildeprint.nl

Copyright © 2024 by J. M. F. Marques

All rights reserved. No part of this book may be reproduced, stored in a retrieval system, or transmitted, in any form or by any means, without prior permission from the copyright owner.

Casimir PhD Series, Delft-Leiden

ISBN 978-94-6384-647-9

An electronic version of this dissertation is available at <http://repository.tudelft.nl/>.

Geweldige dingen worden niet gedaan door impuls, maar door een reeks kleine dingen bij elkaar gebracht.

Vincent Van Gogh

CONTENTS

Summary / Samenvatting / Zämefassig	xi
1 Theory of superconducting transmon qubits	1
1.1 Circuit quantum electrodynamics	2
1.1.1 Single-qubit gates	4
1.1.2 Two-qubit gates	5
1.1.3 Readout	6
2 Automated calibration and benchmarking of superconducting qubit devices	9
2.1 Introduction	10
2.2 Single-qubit gate calibration	10
2.3 Two-qubit gate calibration.	12
2.3.1 Flux to detuning conversion	12
2.3.2 Characterizing the 11-02 interaction.	13
2.3.3 Sudden net-zero gate	15
2.4 Dealing with interacting two-level systems	18
2.4.1 Impact of TLSs in single-qubit gates and readout	19
2.4.2 Impact of TLSs in two-qubit gates	20
2.4.3 Unipolar two-qubit gates	22
2.5 Assessing performance and progress	24
3 Logical qubit operations in an error detecting surface code	27
3.1 Introduction	28
3.2 Results	29
3.2.1 Stabilizer measurements	29
3.2.2 Logical state initialization using stabilizer measurements	31
3.2.3 Logical measurement of arbitrary states	31
3.2.4 Logical gates	33
3.2.5 Pipelined versus parallel stabilizer measurements	35
3.3 Discussion.	36
3.4 Methods	38
3.4.1 Device.	38
3.4.2 State tomography	38
3.4.3 Process tomography in the codespace	39

3.4.4	Extraction of error-detection rate	40
3.5	Supplementary Information	40
3.5.1	Device characteristics	40
3.5.2	Parity-check performance	42
3.5.3	Process tomography	42
3.5.4	Logical state stabilization	42
3.5.5	Fault tolerance of logical operations	44
3.5.6	Quantifying the logical assignment fidelity	47
3.5.7	Numerical analysis	49
4	All-microwave leakage reduction units for quantum error correction with superconducting qubits	59
4.1	Introduction	60
4.2	Results	62
4.2.1	All-microwave leakage reduction unit scheme	62
4.2.2	Calibration of the leakage reduction unit pulse	62
4.2.3	Benchmarking in the qubit subspace	64
4.2.4	Reduction of leakage in repeated stabilizer measurements	64
4.3	Discussion	67
4.4	Supplementary information	68
4.4.1	Device	69
4.4.2	Repeated LRU application	70
4.4.3	Estimating effective transition coupling $\tilde{g}^f$	71
4.4.4	Population in the readout resonators	73
4.4.5	Drive crosstalk	74
4.4.6	Benchmarking the weight-2 parity check	75
4.4.7	Measurement-induced transitions	75
4.4.8	Readout of transmon states	78
4.4.9	Leakage reduction for higher states	78
5	Benchmarking for quantum error correction experiments	83
5.1	Introduction	84
5.2	Benchmarking quantum parity checks	85
5.2.1	Parity assignment fidelity	85
5.2.2	Projection onto stabilizer subspace	86
5.2.3	Repeatability	87
5.3	Signature of errors in repeated stabilizer measurements	88
5.3.1	Ancilla qubit and measurement errors	90
5.3.2	Leakage errors	92
5.3.3	Data qubit errors	94

5.4	Correcting bit flips in a distance 3 surface code	95
5.4.1	Assessing physical qubit errors <i>in situ</i>	95
5.4.2	Logical error rate	98
6	Outlook and conclusion	101
6.1	Outlook	102
6.2	Looking into the future	104
A	Appendix A	107
A.1	Analytical model of SNZ landscape.	108
	Acknowledgements	111
	Curriculum Vitæ	115
	List of Publications	117
	References	119

SUMMARY

The advent of the computer has ushered in the fastest period of technological progress experienced by civilization. Quantum computing leverages the principles of quantum mechanics to process information in a fundamentally different paradigm, one that enables more efficient solving of computationally complex problems, including the factoring of large numbers, optimizing complex systems, and simulating molecular structures. In this new paradigm, information is encoded in quantum bits (qubits), which inherit the wave-like properties of superposition and interference to facilitate more efficient algorithms. However, to realize qubits experimentally, one must encode them into a quantum degree of freedom of a physical system. These are typically hard to control with very high accuracy. Interactions with the environment often lead to errors, limiting the number of operations one can perform reliably in these systems. Quantum error correction provides an alternative to protect quantum information from errors due to decoherence and operational imperfections at the cost of redundancy. Its core idea is to create a highly accurate logical qubit from many noisy physical qubits, ensuring that computational integrity is maintained despite the presence of errors.

This thesis studies experimental aspects of implementing quantum error correction with superconducting qubits: qubits encoded in quantum states of superconducting circuits operating at microwave frequencies and cooled down to cryogenic temperatures where they can exhibit coherent quantum behavior.

In chapter 1, we briefly introduce the theory of circuit quantum electrodynamics. This theory establishes an architecture for quantum computers based on superconducting circuits. We focus particularly on the transmon, a superconducting circuit with a weakly anharmonic spectrum, which has experienced remarkable progress throughout recent years in both scale and quality. We discuss the physical properties of this system and how it can be used to encode a qubit. Additionally, we explain its basic principles of operation for single- and two-qubit gates as well as readout on this platform.

In chapter 2, we discuss the calibration, control, and benchmarking methods used for single- and two-qubit gates and measurement in the experiments reported in this thesis. Calibration is one of the most time-consuming steps of running a superconducting processor. Therefore, when working with large multi-qubit processors, it becomes imperative to develop automated and efficient calibration protocols. We describe the calibration process of single- and two-qubit gates on our platform with emphasis on accuracy and runtime. By employing physical models to analyze calibration sweeps and extract quantities of interest, we show that we can achieve high calibration accuracy while reducing the amount of sampling in param-

eter space, ultimately decreasing the overall calibration runtime. We also investigate errors caused by spurious interactions with strongly coupled two-level systems. These are the leading cause of outliers in gate error, which end up limiting the overall performance of experiments on our devices. We present techniques to mitigate the impact of two-level systems on gate performance by avoiding resonance between the transmon levels and these modes. Finally, we reflect on the experimental progress of subsequent device generations by looking at the distribution of errors of readout, single-, and two-qubit gates.

In chapter 3, we demonstrate the operation of an error-detected logical qubit in the surface code. Future fault-tolerant quantum computers must be capable of storing and processing quantum information encoded in logical qubits. We realize this concept in an error-detected logical qubit encoded in seven physical qubits. A logical precursor to quantum error correction, an error-detected qubit operates using many of the same elements that one also requires for quantum error correction. We detect errors in a distance-two surface code logical qubit—and discard those occurrences—by measuring defects in stabilizer syndromes. Using this, we perform error-detected logical operations including arbitrary initialization, measurement, and gates as well as repeated stabilizer measurements. For each type of operation, we study the performance of fault-tolerant and non-fault-tolerant versions.

In chapter 4, we realize an all-microwave leakage reduction unit for quantum error correction. Standard stabilizer codes are not fault-tolerant to leakage errors. When qubit population leaks outside the computational subspace, it can remain there for many error correction cycles, causing multiple errors over time. Additionally, leakage can spread to other neighboring qubits during two-qubit gate operations. The former and latter lead to correlated errors which are particularly detrimental to logical error rates. To mitigate this issue, leakage reduction units can be interleaved with error correction cycles to bring leaked population back to the computational subspace. Doing so converts a leakage event into a Pauli-like error. We present a scheme with minimal requirements to implement such an operation using a microwave-frequency pulse. We show that this operation simultaneously removes leakage population effectively while having minimal impact on the qubit subspace. As a first application in the context of quantum error correction, we then show how multiple simultaneous LRUs can reduce the error detection rate and suppress leakage buildup within 1% in data and ancilla qubits over 50 cycles of a weight-2 stabilizer measurement.

In chapter 5, we study benchmarking techniques for quantum error correction experiments. Stabilizer codes rely on measuring stabilizer observables using quantum parity checks. These are designed to discretize, detect, and correct errors without disturbing the encoded logical quantum state. The outcome of stabilizer measurements is used to identify error syndromes that signal the occurrence of various types of qubit errors. By understanding how physical error mechanisms manifest within stabilizer syndromes, we can quantify distinct error signatures ranging from simple bit- or phase-flips to more complex leakage errors. Finally, we perform repeated error correction of bit-flip errors in a distance-3 surface code.

SAMENVATTING

De komst van de computer heeft de snelste periode van technologische vooruitgang in de geschiedenis van de beschaving ingeluid. Quantum computing maakt gebruik van de principes van de quantummechanica om informatie te verwerken in een fundamenteel ander paradigma, dat het efficiënter oplossen van computationeel complexe problemen mogelijk maakt, zoals het factoriseren van grote getallen, het optimaliseren van complexe systemen en het simuleren van moleculaire structuren. In dit nieuwe paradigma wordt informatie gecodeerd in quantum bits (qubits), die de golfachtige eigenschappen van superpositie en interferentie erven om efficiëntere algoritmen mogelijk te maken. Echter, om qubits experimenteel te realiseren, moet men ze coderen in een quantumvrijheidsgraad van een fysiek systeem. Deze zijn doorgaans moeilijk met zeer hoge nauwkeurigheid te controleren. Interacties met de omgeving leiden vaak tot fouten, waardoor het aantal operaties dat men betrouwbaar kan uitvoeren in deze systemen beperkt wordt. Quantumfoutcorrectie biedt een alternatief om quantuminformatie te beschermen tegen fouten als gevolg van decoherentie en operationele imperfecties, tegen de prijs van overtolligheid. Het is om een zeer nauwkeurige logische qubit te creëren uit veel fysieke qubits, zodat de computationele integriteit behouden blijft ondanks de aanwezigheid van fouten.

Dit proefschrift bestudeert experimentele aspecten van het implementeren van quantumfoutcorrectie met supergeleidende qubits: qubits gecodeerd in quantumtoestanden van supergeleidende circuits die werken op microgolf-frequenties en afgekoeld worden tot cryogene temperaturen waar ze coherent quantumgedrag kunnen vertonen.

In hoofdstuk 1 introduceren we kort de theorie van circuit quantum electrodynamics. Deze theorie vestigt een architectuur voor quantumcomputers op basis van supergeleidende circuits. We richten ons vooral op de transmon, een supergeleidende schakeling met een zwak anharmonisch spectrum, die de afgelopen jaren opmerkelijke vooruitgang heeft geboekt zowel in schaal als kwaliteit. We bespreken de fysische eigenschappen van dit systeem en hoe het gebruikt kan worden om een qubit te coderen. Daarnaast leggen we de basisprincipes uit van de werking van enkel- en tweequbit, evenals op dit platform.

In hoofdstuk 2 bespreken we de kalibratie-, controle- en benchmarkingmethoden die gebruikt worden voor enkel- en tweequbit en meting in de experimenten die in dit proefschrift worden gerapporteerd. Kalibratie is een van de tijdrovendste stappen bij het gebruik van een supergeleidende processor. Daarom wordt het bij het werken met grote multiqubit-processors noodzakelijk om geautomatiseerde en efficiënte kalibratieprotocollen te ontwikkelen. We beschrijven het kalibratieproces van enkel- en tweequbit op ons platform met de

nadruk op nauwkeurigheid en doorlooptijd. Door fysische modellen te gebruiken om kalibratieafspraken te analyseren en hoeveelheden van belang te extraheren, laten we zien dat we een hoge kalibratienauwkeurigheid kunnen bereiken terwijl we de hoeveelheid sampling in de parameter ruimte verminderen, wat uiteindelijk de totale kalibratietijd vermindert. We onderzoeken ook fouten veroorzaakt door spurious interacties met sterk gekoppelde twee-niveausystemen. Deze zijn de belangrijkste oorzaak van uitschieters in fouten, die uiteindelijk de algehele prestaties van experimenten op onze apparaten beperken. We presenteren technieken om de impact van twee-niveausystemen op de prestaties te verminderen door resonantie tussen de transmon-niveaus en deze te vermijden. Tot slot reflecteren we op de experimentele vooruitgang van opeenvolgende apparaatgeneraties door te kijken naar de verdeling van fouten bij uitlezing, enkel- en tweequbit.

In hoofdstuk 3 demonstreren we de werking van een foutgedetecteerde logische qubit in de oppervlaktecode. Toekomstige fouttolerante quantumcomputers moeten in staat zijn om quantuminformatie op te slaan en te verwerken die is gecodeerd in logische qubits. We realiseren dit concept in een foutgedetecteerde logische qubit gecodeerd in zeven fysieke qubits. Een logische voorloper van quantumfoutcorrectie, een foutgedetecteerde qubit werkt met veel van dezelfde elementen die ook nodig zijn voor quantumfoutcorrectie. We detecteren fouten in een logische qubit met afstand twee oppervlaktecode—en verwerpen die gevallen—door defecten in stabilizersyndromen te meten. Met behulp hiervan voeren we foutgedetecteerde logische operaties uit, waaronder willekeurige initiatie, meting en poorten, evenals herhaalde stabilizermetingen. Voor elk type operatie bestuderen we de prestaties van fouttolerante en niet-fouttolerante versies.

In hoofdstuk 4 realiseren we een all-microwave leakage reduction unit voor quantumfoutcorrectie. Standaard stabilizer-codes zijn niet fouttolerant voor leakage fouten. Wanneer qubitpopulatie buiten de computationele subruimte lekt, kan het daar vele foutcorrectiecycli blijven, wat leidt tot meerdere fouten in de tijd. Bovendien kan leakage zich verspreiden naar andere naburige qubits tijdens tweequbitpoort-operaties. Het eerste en het laatste leiden tot gecorreleerde fouten die bijzonder schadelijk zijn voor logische foutpercentages. Om dit probleem te verminderen, kunnen leakage reduction units worden ingeschakeld tussen foutcorrectiecycli om gelekte populatie terug te brengen naar de computationele subruimte. Dit converteert een leakage-gebeurtenis in een Pauli-achtige fout. We presenteren een schema met minimale vereisten om een dergelijke operatie te implementeren met behulp van een microgolf-frequentie puls. We laten zien dat deze operatie tegelijkertijd leakage-populatie effectief verwijdert terwijl het minimale impact heeft op de qubitsubruimte. Als eerste toepassing in de context van quantumfoutcorrectie laten we vervolgens zien hoe meerdere gelijktijdige LRUs de foutdetectiesnelheid kunnen verminderen en leakage-opbouw kunnen onderdrukken tot minder dan 1% in data- en ancillaqubits gedurende 50 cycli van een gewicht-2 stabilizermeting.

In hoofdstuk 5 bestuderen we benchmarkingtechnieken voor quantumfoutcorrectie-experimenten. Stabilizer-codes zijn afhankelijk van het meten van stabilizer-observabelen met behulp van quantumpariteitscontroles. Deze zijn ontworpen om fouten te discretiseren, detecteren en corrigeren zonder de gecodeerde logische quantumtoestand te verstoren. De uitkomst van stabilizermetingen wordt gebruikt om fouten-syndromen te identificeren die het optreden van verschillende soorten qubitfouten signaleren. Door te begrijpen hoe fysische foutmechanismen zich manifesteren binnen stabilizersyndromen, kunnen we verschillende foutsignatures kwantificeren, variërend van eenvoudige bit- of fasefouten tot meer complexe leakage-fouten. Tot slot voeren we herhaalde foutcorrectie uit van bitflip-fouten in een oppervlaktecode met afstand 3.

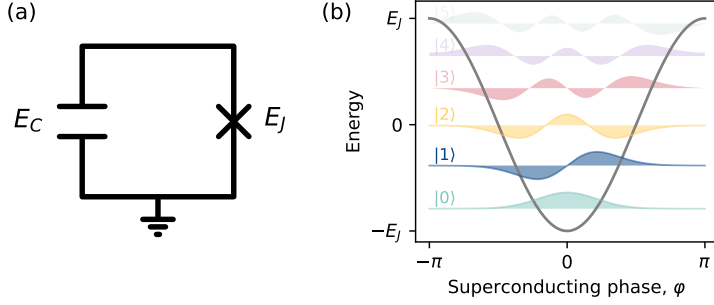


Figure 1.1: **Transmon circuit and its wavefunctions.** (a) Transmon circuit consisting of a Josephson junction (right) shunted by a capacitor (left). (b) Wavefunctions, $\Re[\psi_i(\varphi)]$, of the transmon plotted within the Josephson potential (grey). Each wave function $\psi_i(\varphi) \equiv |i\rangle$ is plotted at its eigenvalue coordinate.

1.1 Circuit quantum electrodynamics

SUPERCONDUCTING circuits use superconducting elements cooled to millikelvin temperatures and operated at microwave frequencies where coherent quantum phenomena can emerge. The study of its behavior in this quantized regime has been termed circuit quantum electrodynamics [1, 2]. Perhaps the most popular instance of such circuits is the transmon [3]. It consists of a Josephson junction [4] shunted by a capacitance as shown in Fig. 1.1a. In a Josephson junction, two superconductors are separated by a thin insulating barrier where tunneling of charge carriers can occur. This gives rise to a potential difference across the junction

$$V = -E_J \cos \varphi, \quad (1.1)$$

where φ is phase difference between the two superconducting states on each side of the junction [5] and $EJ = \Phi_0 I_C / 2\pi$ is the junction specific Josephson energy which depends on the critical current it supports, I_C , and the magnetic flux quantum Φ_0 . The near-lack of dissipation in these elements allows the circuit to exhibit quantum coherence and the non-linearity conferred by the Josephson junction enables its manipulation for quantum information processing. Although superconducting circuits possess many microscopic degrees of freedom, their dynamics can be described by just a few macroscopic ones. To this end, rigorous quantization techniques have been developed to quantize the dynamics of these circuits [6]. In the particular case of the transmon [3], two observables describe its macroscopic behavior: the charge difference between the two sides of the capacitor, $\hat{Q}$, and the superconducting phase difference across the junction, $\hat{\varphi}$. The Hamiltonian associated to this system is given by

$$\hat{H} = \frac{(\hat{Q} - Q_g)^2}{2C} - E_J \cos \hat{\varphi}, \quad (1.2)$$

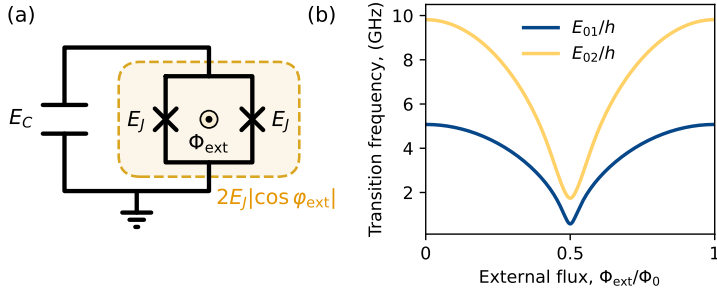


Figure 1.2: **Flux-tunable transmon.** (a) Flux-tunable transmon using a superconducting quantum interference device (SQUID) loop (yellow). (b) Transition frequency between the ground state $|0\rangle$ and the first $|1\rangle$ and second $|2\rangle$ excited states of the transmon as a function of external flux bias. Transitions to higher excited states are denoted in grey.

where the first and second terms correspond to the energy stored in the capacitance, C , relative to a charge offset, Q_g , and the Josephson potential, respectively. It is common express the former term using instead the charge number observable $\hat{n} = \hat{Q}/2e$. This leads to the commonly known transmon Hamiltonian [3],

$$\hat{H} = 4E_C(\hat{n} - n_g)^2 - E_J \cos \hat{\varphi}, \quad (1.3)$$

where $E_C = e^2/2C$, $n_g = Q_g/2e$ and whose conjugate observables $\hat{n}$ and $\hat{\varphi}$ satisfy the commutation rule

$$[\hat{n}, \hat{\varphi}] = -i. \quad (1.4)$$

The eigenstates of Eq. 1.3 can be better understood in the phase basis, $\hat{\varphi}$. In this basis, the hamiltonian becomes analogous to that of a mass suspended in a pendulum [3]: the charging energy corresponding to the kinetic energy of the mass and the Josephson potential corresponding to the gravitational energy. Thus we can think of charge $\hat{n} = -i \frac{\partial}{\partial \varphi}$ and $\cos \varphi$ as the canonical momentum and position observables in the pendulum, respectively. In the transmon regime where $E_J \gg E_C$ [3], the lowest energy states of the system are bounded within the cosine potential as shown Fig. 1.1b. In the pendulum analogy this corresponds to small oscillations of the mass around the equilibrium point. Conversely, in the transmon, this corresponds to oscillations of charge and current between the two branches of the circuit. Most importantly, the spectrum of eigenvalues of these states is anharmonic (non-linear). In fact, one can show that the anharmonicity is directly proportional $-E_C$ in this regime [2]. This allows addressing individual transitions to control the dynamics between eigenstates coherently. Furthermore, this regime suppresses the eigenvalue dependence on the charge offset n_g [3] (known as charge dispersion) of states within the cosine potential. This allows us to neglect charge offsets in the system which are generally hard to control and would otherwise constitute a source of noise. Instead, it is common to introduce a second junction

to form a DC-SQUID [5] (superconducting quantum interference device) as shown in Fig. 1.2a. This allows us to tune the effective Josephson energy of the branch using a magnetic flux. When subject to a flux bias, Φ_{ext} , flux quantization results in the condition

$$\hat{\varphi}_1 - \hat{\varphi}_2 + 2\varphi_{\text{ext}} = 0 \pmod{2\pi}, \quad (1.5)$$

where $\varphi_{\text{ext}} = \frac{\pi\Phi_{\text{ext}}}{\Phi_0}$. The effective Josephson potential is therefore given by,

$$-E_J \cos \hat{\varphi}_1 - E_J \cos \hat{\varphi}_2 = -2E_J |\cos \varphi_{\text{ext}}| \cos \hat{\varphi} \quad (1.6)$$

where $\hat{\varphi} = (\hat{\varphi}_1 + \hat{\varphi}_2)/2$. Figure 1.2b shows the transition frequencies E_{01}/h and E_{02}/h as function of external flux. Flux control will be key for many two-qubit gate architectures [7–10], including the one used in this thesis [11].

Circuit QED architecture [1] uses the lowest two energy states of the transmon, $|0\rangle$ and $|1\rangle$, to encode qubits in digital quantum computers. Operations on qubits are the remaining requirements needed to operate such quantum computers [12]. In the next subsection we discuss how to perform single- and two-qubit gates and measurement in this platform.

1.1.1 Single-qubit gates

To operate single-qubit gates on the transmon, it is common to use microwave drives resonant with the qubit transition. To understand the effect of a drive on the transmon, it is useful to introduce the E_J , E_C harmonic oscillator operators $\hat{b}$, $\hat{b}^\dagger$. This basis - which diagonalizes the hamiltonian $H_{\text{HO}} = 4E_C \hat{n}^2 + \frac{1}{2}E_J \hat{\varphi}^2$ - is defined as:

$$\hat{\varphi} = \left(\frac{2E_C}{E_J} \right)^{1/4} (\hat{b}^\dagger + \hat{b}), \quad (1.7)$$

$$\hat{n} = i \left(\frac{E_J}{32E_C} \right)^{1/4} (\hat{b}^\dagger - \hat{b}). \quad (1.8)$$

In the transmon regime, the matrix elements of $\hat{b}$ and $\hat{b}^\dagger$, in the transmon eigenbasis, resemble bosonic annihilation and creation operators for the first eigenstates of the transmon circuit (i.e., eigenstates bounded within the cosine potential well). We can then write the drive hamiltonian as,

$$H_d(t) = \frac{\hbar\Omega}{n_{\text{ZPF}}} V_d(t) \hat{n}. \quad (1.9)$$

where, $n_{\text{ZPF}} = \left(\frac{E_J}{32E_C} \right)^{1/4}$ denotes the charge vacuum state fluctuations of H_{HO} , $V_d(t)$ is the time-dependent drive seen by the transmon, and Ω the coupling of the drive. The full system hamiltonian, truncated down to the first two eigenstates, can be written as [13]

$$H/\hbar = -\frac{\omega_Q}{2} \hat{\sigma}_z + \Omega V_d(t) \hat{\sigma}_y. \quad (1.10)$$

where $\hat{\sigma}_z$ and $\hat{\sigma}_y$ are Pauli matrices. Under the bare qubit hamiltonian (first term in Eq. 1.10), the qubit eigenstates evolve around the Z -axis of the Bloch sphere with angular frequency

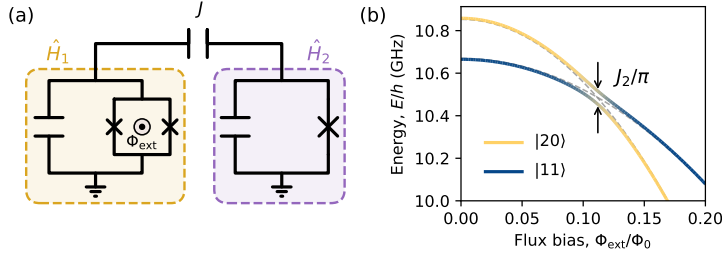


Figure 1.3: **Transmon-transmon coupling.** (a) A flux-tunable transmon (yellow) capacitively coupled to a fixed-frequency transmon (purple). (b) Eigenvalues of the composite system states versus external flux bias Φ_{ext} . States are colored according to their overlap with the bare states $|20\rangle$ and $|11\rangle$ (denoted by the dashed curves).

ω_Q . By moving to frame rotating at the qubit frequency, the first term of H vanishes and the remaining drive term turns into [13],

$$\tilde{H}/\hbar = \Omega V_d(t) [\cos(\omega_Q t) \hat{\sigma}_x + \sin(\omega_Q t) \hat{\sigma}_y]. \quad (1.11)$$

To better understand the behavior of the qubit in this frame, we assume a general drive expression $V_d(t) = s(t)(I \sin(\omega_D t) + Q \cos(\omega_D t))$ characterized by a pulse envelope $s(t)$, frequency ω_D and *in-phase* and *out-of-phase* components I and Q . On resonance i.e., $\omega_D = \omega_Q$, and after dropping fast rotating terms which have negligible effect on the qubit dynamics [2], the driven hamiltonian takes the simple form:

$$\tilde{H}/\hbar = \frac{s(t)\Omega}{2} [I\sigma_x + Q\hat{\sigma}_y]. \quad (1.12)$$

Therefore, in the rotating frame, the qubit will evolve under $\hat{\sigma}_x$ or $\hat{\sigma}_y$ for a drive along the I or Q quadrature, respectively. This means that one can perform arbitrary single-qubit rotations around X and Y by changing the phase of the drive. Additionally, one can perform rotations around Z by simply changing the rotating frame. This is known as virtual- Z gate, and is performed by incrementing the phase of subsequent drives by the desired rotation angle.

1.1.2 Two-qubit gates

In this subsection, we present the two-qubit gate architecture used in the reported experiments. A common method for coupling superconducting circuits uses direct capacitive coupling (see Fig.1.3a). Since the resulting circuit is now a two node network, the capacitive energy of the circuit is a function of charge and voltage on all nodes [13]:

$$\frac{1}{2} \vec{Q} \cdot \vec{V} = \frac{1}{2} \vec{Q}^T \cdot \mathbf{C}^{-1} \cdot \vec{Q} \quad (1.13)$$

where $\vec{Q}$ and $\vec{V}$ are vectors denoting charge and voltage on each node and $\mathbf{C}$ is the capacitance matrix [14] of the network. One can see from Eq. 1.13, that this will result in an interaction Hamiltonian of the form:

$$H_{\text{coupling}}/\hbar = -J\hat{n}_1\hat{n}_2, \quad (1.14)$$

where $\hat{n}_1$ and $\hat{n}_2$ represent the charge number operators for each transmon. In fixed-coupling architectures, transmons are tuned in and out of resonance to effectively turn on and off the interaction [7]. Specifically, one must ensure that the detuning between the relevant levels is significantly larger than the coupling when the transmons are idling. For this purpose, we use a high-frequency flux-tunable transmon, which is biased into resonance with a lower frequency transmon, as depicted in Figure 1.3b. For our two-qubit gate implementation, we leverage the interaction between the computational state $|11\rangle$ and the non-computational state $|20\rangle$. Here, we follow the naming convention $|ij\rangle$ where i and j denote the state of transmon 1 and 2, respectively. When these states are resonant, the coupling leads to an avoided crossing of magnitude:

$$2\hbar J_2 = 2|\langle 20|H_{\text{coupling}}|11\rangle|, \quad (1.15)$$

hybridizing the two states as shown in Fig. 1.3b. This interaction can be exploited to implement a controlled phase gate in two ways. By moving in and out of the avoided crossing adiabatically [15], such that the system remains in its eigenstate, the repulsion $\hbar\zeta = (\tilde{E}_{11} - \tilde{E}_{10} - \tilde{E}_{01} + \tilde{E}_{00})$ caused by the interaction leads to relative phase accrual of state $|11\rangle$ relative to other states in the computational subspace. Here, $\tilde{E}_{ij}$ denotes the energy of the dressed state that most overlaps with the bare state $|ij\rangle$ at zero flux bias. The total relative phase acquired during the flux trajectory $\phi_c = \int \zeta(t)dt$ (modulo single-qubit phases also accrued in the process) can be used to perform a conditional phase (CPhase) gate of the form

$$U_{\text{CPhase}} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & e^{i\phi_c} \end{pmatrix}, \quad (1.16)$$

By instead moving diabatically into the interaction, one can leverage exchange between the two states to implement a Cphase gate [11, 16, 17]. To understand this, one can think of how the state $|11\rangle$ evolves on resonance. After a period $\tau = \pi/J_2$ the state performs a full rotation in the $\{|11\rangle, |20\rangle\}$ manifold, acquiring a phase $\phi_c = \pi$ in the process. In fact, as shown in Figure 2.13, this procedure can also be generalized to realize an arbitrary CPhase.

1.1.3 Readout

In circuit QED, readout of the state of a superconducting circuit is performed by means of an auxiliary resonator mode. This is usually achieved by coupling the transmon node to a

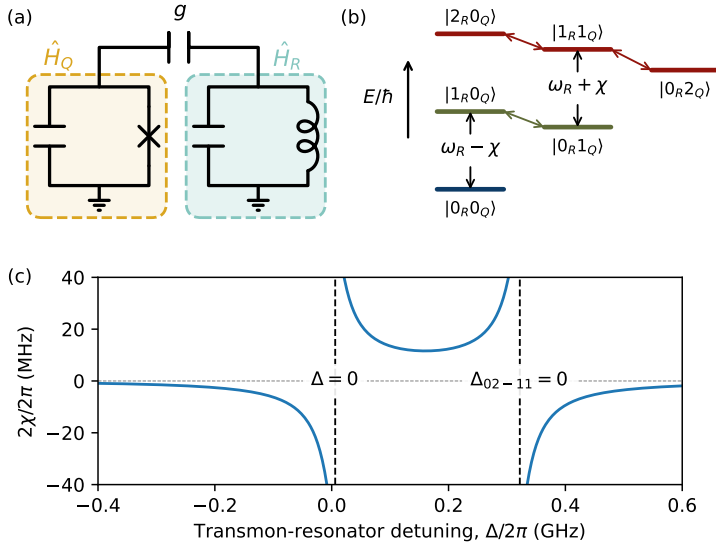


Figure 1.4: **Dispersive readout in the transmon.** (a) Transmon circuit (H_Q) capacitively coupled with coupling strength g to an LC oscillator (H_R). (b) Energy level ladder of the joint transmon and oscillator system where arrows denote non-zero coupling between levels. The dispersive shift 2χ can be understood by the difference between the two resonator frequencies when the qubit is in $|0\rangle$ ($\omega_R - \chi$) and in $|1\rangle$ ($\omega_R + \chi$). (c) Dispersive shift as function of qubit-resonator detuning $\Delta = \omega_Q - \omega_R$. The dispersive shift is maximized when the two systems are resonant ($\Delta = 0$) or when the resonator is resonant the $|1\rangle \rightarrow |2\rangle$ transition of the transmon ($\Delta_{02-11} = 0$).

terminated transmission line (see for eg. Fig. 3.1a). Although such elements have many resonance modes (characterized by the input impedance seen by the transmon circuit [14]), we will restrict our analysis to a single resonance mode modeled by a lumped element LC oscillator as shown in Fig. 1.4a. Analogously to the previous subsection, the coupling capacitance between the transmon and the LC oscillator will result in a coupling term described by the Eq. 1.14. The full system Hamiltonian of the coupled oscillator system is thus described by,

$$H/\hbar = H_T/\hbar + \omega_R \left(\hat{a}^\dagger \hat{a} + \frac{1}{2} \right) - g(\hat{a}^\dagger - \hat{a})(\hat{b}^\dagger - \hat{b}). \quad (1.17)$$

The first term in Eq. 1.17 denotes the bare transmon Hamiltonian, the second term the bare resonator hamiltonian - described by a harmonic oscillator at angular frequency ω_R and bosonic operators $\hat{a}$ and $\hat{a}^\dagger$ - and the final term the coupling between both systems. Note that, the coupling constant, g , has now absorbed the charge vacuum fluctuations. We follow by dropping terms proportional to $\hat{a}\hat{b}$ and $\hat{a}^\dagger\hat{b}^\dagger$. This approximation, known as the rotating-wave-approximation (RWA), is valid here under the assumption that such terms will

1

only couple far detuned levels thus having a negligible effect. This holds true when the detuning between the qubit and the resonator, $\Delta = \omega_Q - \omega_R$, is much lower than their transition frequencies ($\Delta \ll \omega_Q, \omega_R$) and if the dynamics of the system is confined to low photon values in the qubit and resonator. We note that this approximation is often broken in experiments when driving the resonator with higher amplitudes [18, 19] (see for eg. Chapter 4) which can lead to leakage of transmon population to higher excited states [20]. After the RWA and truncating the transmon to a two-level system, the resulting system is described by the Jaynes-Cummings Hamiltonian [2]:

$$H/\hbar = -\frac{\omega_Q}{2} \hat{\sigma}_Z + \omega_R \left(\hat{a}^\dagger \hat{a} + \frac{1}{2} \right) + g(\hat{a}^\dagger \hat{\sigma}_- + \hat{a} \hat{\sigma}_+) \quad (1.18)$$

In the dispersive regime, where the qubit-resonator detuning is significantly larger than the coupling, $\Delta \gg g$, one can show that, using second-order perturbation theory [1], the Hamiltonian can be approximated by:

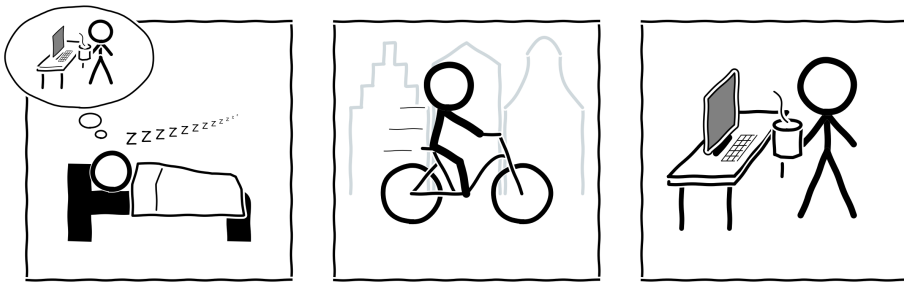
$$H/\hbar = (\omega_R - \chi \hat{\sigma}_Z) \left(\hat{a}^\dagger \hat{a} + \frac{1}{2} \right) - \frac{\tilde{\omega}_Q}{2} \hat{\sigma}_Z \quad (1.19)$$

where $\chi = g^2/\Delta$ is known as the dispersive shift. Its physical meaning can be interpreted as a qubit-state dependent shift of the resonator frequency. Additionally, $\tilde{\omega}_Q = \omega_Q + g^2/\Delta$ is now the Lamb-shifted frequency of the qubit. In circuit QED, the dispersive shift χ is leveraged for quantum non-demolition readout of the qubit state. This is executed by probing the scattering properties of the resonator (transmission or reflection) effectively performing a projective measurement of the qubit system. Higher levels present in the transmon (as shown in Fig. 1.4b) modify the dispersive shift expression to:

$$\chi = \frac{g^2}{\Delta} \left(\frac{1}{1 - \Delta/E_C} \right). \quad (1.20)$$

The outcome of this is plotted in Fig. 1.4c. This shows, how one can engineer the coupling and qubit/resonator frequencies, g and Δ , for optimal readout parameters. In the following chapters, we'll discuss measurement and tune-up protocols to calibrate these operations effectively.

AUTOMATED CALIBRATION AND BENCHMARKING OF SUPERCONDUCTING QUBIT DEVICES



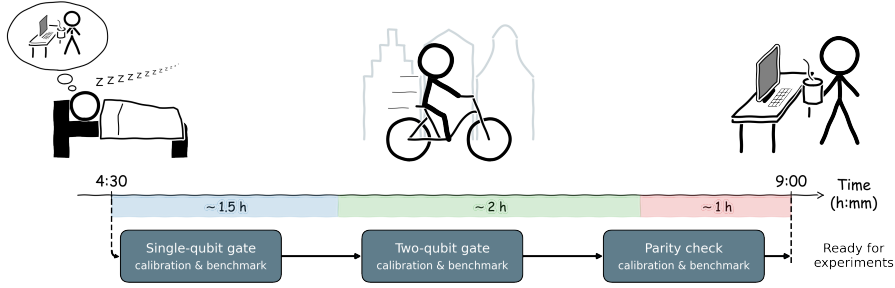


Figure 2.1: **Automated calibration procedure.** In order to make the most out of daily experimental use, the calibration procedure (which has a duration of about 4.5 hours) is initiated at 4:30 AM and concludes at 9:00 AM. It performs sequential calibration and benchmark of single-qubit, two-qubit and parity check operations.

2.1 Introduction

The operation of large-scale quantum processors requires accurate, reliable, and fast automated calibration protocols. Within the framework of a Surface-17 device, comprising 17 single-qubit gates and 24 two-qubit gates, all susceptible to parameter drift, the calibration process becomes inherently time-consuming due to the involvement of numerous repetitive tasks. Hence, to streamline and enhance the efficiency of the calibration process for daily use, automation becomes imperative. Moreover, daily benchmarking of the processor is essential to ascertain that its performance aligns with the stringent requirements of our experiment. An automated calibration framework, integral to this process, should also establish an abstraction layer facilitating the exploration of various parameter regimes. In particular, explore different configurations of idle and interaction frequencies. This abstraction layer was essential for mitigating the impact of spurious TLSs (and their associated drift) on the quantum device, ensuring the robustness and reliability of the processor's performance.

In this chapter, I describe the automated calibration framework developed for a Surface-17 device featuring the pipelined architecture [21]. Its procedure, illustrated conceptually in Fig. Figure 2.1, spans approximately 4.5 hours. It is initiated at 4:30 AM, strategically planned to conclude around 9:00 AM. The calibration protocol includes three steps, first involving the calibration and benchmarking of single-qubit and two-qubit gates and metrics. Subsequently, a crucial third step focuses on parity-check calibration and benchmarking, tailored for the quantum error correction experiments performed in this thesis.

2.2 Single-qubit gate calibration

Single-qubit gates are performed using a microwave pulse with DRAG [23]. A microwave drive allows for rotation around any axis lying along the XY plane of the Bloch sphere [13].

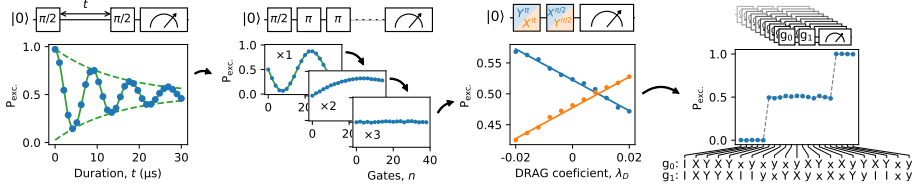


Figure 2.2: **Single-qubit gate calibration procedure.** The circuit (top) and measurement data (bottom) is shown for each step of the calibration. The frequency of the pulse is calibrated using a Ramsey experiment to measure the detuning between the drive and the qubit frequency. The drive amplitude is calibrated by using a varying sequence of repeating pi-pulses. This step is repeated iteratively until the measured error is negligible. Following this, we tune the DRAG coefficient, λ_D , such that the measured P_{exc} of the two circuits (blue and orange) is the same. Finally, we assess any remaining coherent errors using an AllXY sequence [22]. This method measures P_{exc} for any two combinations from the set of gates formed by I , $\pi/2$ - and π -pulses around the X and Y axis. Here, coherent errors translate into deviations from the ideal outcome (dashed line) and can be often linked to a common error mechanism [22].

This parametrization includes four parameters: pulse duration (σ_D), frequency (f_D), amplitude (A_D) and DRAG coefficient (λ_D) out of which three are calibrated on a daily basis (f_D , A_D and λ_D). The single-qubit gate calibration procedure used is illustrated in Fig. Figure 2.2. We first use a Ramsey measurement to accurately find the pulse frequency, f_D . Following this, we find the optimal pulse amplitude, A_D , by using an increasingly long sequence of π -pulses designed to amplify small rotation errors in the gate. By fitting the data one can calculate the associated amplitude error of the gate. Due to non-linearity in the RF up-conversion setup used to synthesize the pulse, we perform this process iteratively until the extracted error meets a predefined threshold. We then calibrate the DRAG coefficient, λ_D using a standard technique [22]. Finally, an AllXY [22] measurement is performed to verify whether significant coherent errors occur in the gate. We then benchmark the gate using Clifford randomized benchmarking [24–27] to assess the average single-qubit gate error. Additional benchmarks including measurement of relaxation and coherence time, T_1 and T_2^{cho} , respectively. These are relevant to monitor and understand possible fluctuations in gate performance. Finally, we also benchmark single-shot readout assignment error. All metrics and measured data are displayed in a performance monitor, shown in Figure Figure 2.3a, which grants a convenient overview of single-qubit performance. Runtime for each routine is shown in Figure Figure 2.3b. The entire calibration and benchmark procedure takes about 6 minutes per qubit totaling 1 hour and 30 minutes for the full device.

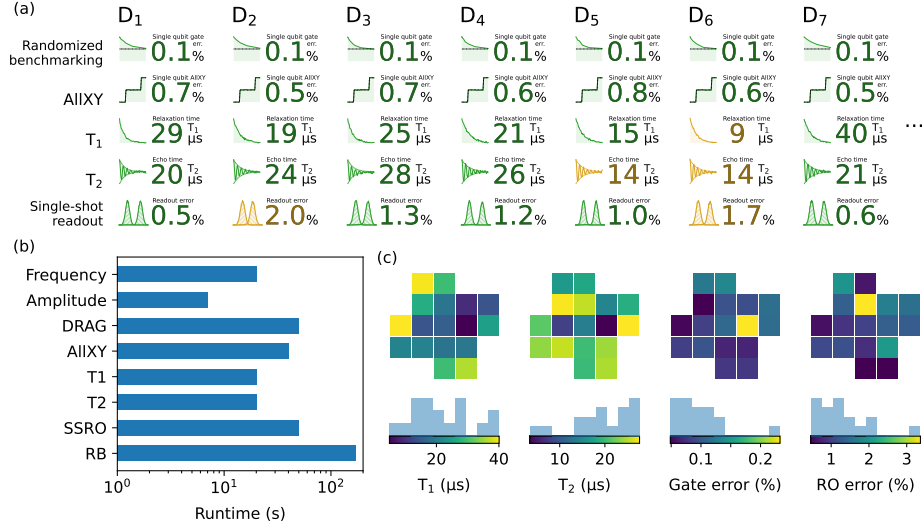


Figure 2.3: **Single-qubit performance monitor and runtime of calibration and benchmark steps.** (a) Relevant single-qubit benchmarks are displayed in the performance monitor for quick visual inspection. Each row displays the relevant metric and (simplified) measurement data from the experiment. (b) Runtime of each calibration and benchmark experiment. (c) Relaxation time, coherence time, single-qubit gate error and single-shot readout assignment error for all qubits in a surface-17 device.

2.3 Two-qubit gate calibration

In our flux-tunable transmon architecture [21], two-qubit gates are executed by dynamically tuning the frequencies of transmons, bringing the joint system levels $|11\rangle$ and $|02\rangle$ in and out of resonance. The calibration framework leverages physical models of this interaction, to efficiently tune gate parameters. In the following subsections, I delineate the steps and models embedded within this procedure.

2.3.1 Flux to detuning conversion

Expressing any physical model of the two-qubit system is most conveniently achieved as a function of transmon detuning. Therefore, our approach relies on the accurate conversion between flux pulse amplitude and transmon detuning. One can measure the detuning experienced by the transmon, $\Delta(A)$, during a square flux pulse of amplitude A by probing the phase of the qubit's Bloch vector as function of pulse duration t using the two Ramsey sequences shown in Fig. Figure 2.4a. From the measured results (Fig. Figure 2.4b) we can find the detuning by calculating the Fast Fourier Transform (FFT) of the complex signal $\langle X \rangle + i\langle Y \rangle$ (Fig. Figure 2.4b). The maximum probable detuning, constrained by the Nyquist

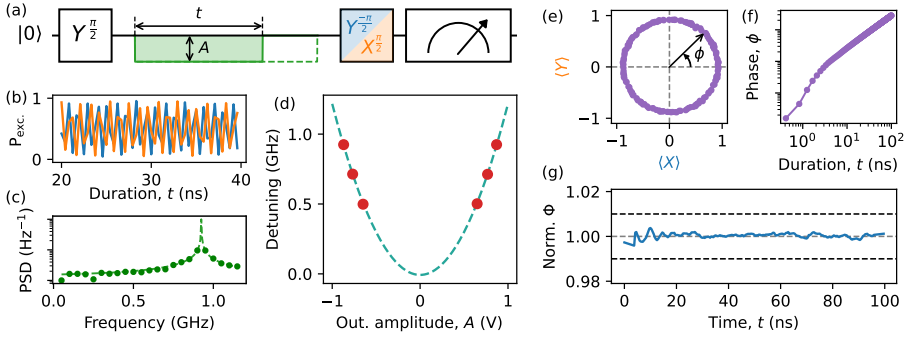


Figure 2.4: Flux to detuning conversion and distortion. (a) Ramsey experiment used to measure the $\langle X \rangle$ (blue) and $\langle Y \rangle$ projection of the Bloch vector after a square flux-pulse (green) of amplitude A and duration t . (b) Raw measurement output as function pulse duration. (c) Fast-Fourier transform of the complex signal, $s(t) = \langle X \rangle + i\langle Y \rangle$. The inferred detuning can be computed by fitting the data to a Lorentzian (dashed curve). (d) Measured detuning for different flux amplitude values and corresponding flux-arc fit (dashed curve). (e-g) Cryoscope procedure [28] for assessing flux distortions. (e) $\langle X \rangle$ and $\langle Y \rangle$ projection for varying flux pulse duration and corresponding Bloch vector phase (f), ϕ , as function of pulse duration. (g) Inferred step response from measured data.

frequency of this measurement, is 1.2GHz (determined by the arbitrary waveform generator (AWG) sampling rate of 2.4GHz). The conventional precision of the FFT would be limited by the total duration of the time trace. To minimize the duration of this measurement, we perform traces of 20ns, resulting in a low precision of 50MHz. However, we address this limitation by fitting the FFT data to a Lorentzian (dashed curve in Fig. Figure 2.4c), achieving a significantly higher precision, typically $\sim 1\text{MHz}$. By measuring the detuning for different flux pulse amplitudes (Fig. Figure 2.4d), we fit a second-order polynomial which describes $\Delta(A)$. Following this, we characterize and calibrate for dynamical pulse distortions using the cryoscope technique [28]. Similar to the previous procedure, this method using a sequence of Ramsey measurements to infer the on-chip step response of the flux pulse (Fig. Figure 2.4e-g). The details of this procedure are described in Ref. [28].

2.3.2 Characterizing the 11-02 interaction

In the next step of the calibration procedure, we characterize the parameters of the $|11\rangle \leftrightarrow |02\rangle$ avoided crossing interaction. The schematic representation of these levels, based on typical parameters, is depicted in Fig. Figure 2.5b. One can make these levels resonant by

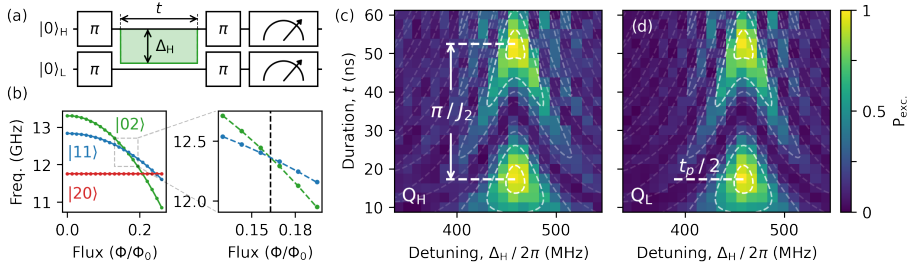


Figure 2.5: Time-domain characterization of $|11\rangle \leftrightarrow |02\rangle$ interaction. (a) Quantum circuit used to characterize the interaction chevron. (b) Energy level diagram for states in the 2-excitation manifold of two uncoupled transmons as function of flux. The exact solution obtained from diagonalizing transmon Hamiltonian using 20 charge states (dots) is used to fit a polynomial for interpolating the intersection of levels. (c) Measured results as function of flux-pulse detuning, Δ_H , and duration, t . The fit of the data to the model of Eq. Equation (2.2) (white dashed contour) allows extracting the relevant parameters without relying on fine sampling of the parameters.

detuning the highest frequency transmon (right index in this notation). A simple model of the system within the $\{|11\rangle, |02\rangle\}$ manifold is described by the Hamiltonian,

$$H = \begin{pmatrix} -\tilde{\Delta}/2 & J_2 \\ J_2 & \tilde{\Delta}/2 \end{pmatrix}, \quad (2.1)$$

where $\tilde{\Delta}$ and J_2 is the detuning and coupling strength, respectively, between states $|11\rangle$ and $|02\rangle$. We solve for the two uncoupled transmon spectra to find the high frequency transmon detuning Δ_0 where the states are resonant, i.e., $\tilde{\Delta} \equiv \Delta_H - \Delta_0 = 0$ (where Δ_H refers to detuning of the high transmon from its idle frequency). To maximize the accuracy of this prediction, we solve for the values E_C and E_J of the high frequency transmon, and obtain the spectrum of $|11\rangle$ and $|02\rangle$ (as function of external Φ) by numerically solving the transmon Hamiltonian in the charge basis (truncating at 15 charge states). We do this for discrete values of Φ and fit a polynomial curve to each state to find the crossing point (Fig. Figure 2.5b). Using this procedure we can correct for flux dependency of the anharmonicity, $\alpha(\Phi)$, which may be relevant for large detunings. Using the sequence depicted in Fig. Figure 2.5a, we sweep the parameters t and Δ_H of a square flux pulse to find an interaction Chevron. The measured landscapes (Fig. Figure 2.5c and Figure 2.5d) are used to fit the model:

$$P_{\text{exc}}(\Delta_H, t, \Delta_0, J_2, \phi_{\text{dist}}) = \left(\frac{2J_2}{\Omega} \right)^2 \left(\frac{1 - \cos(\Omega t - \phi_{\text{dist}})}{2} \right). \quad (2.2)$$

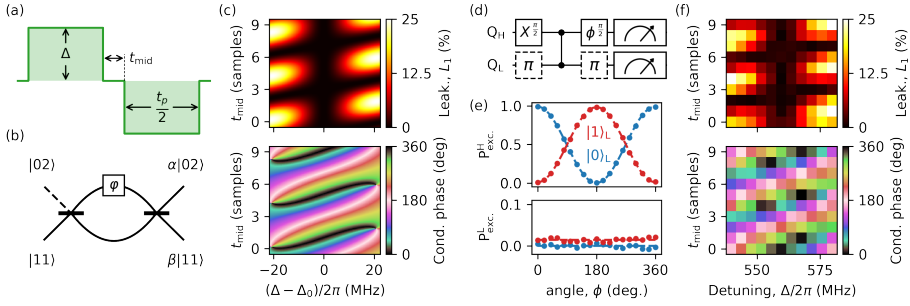


Figure 2.6: Sudden net-zero two-qubit gate. (a) Sudden net-zero gate parametrization. The mechanics of this gate can be better understood by analogy to a Mach-Zender interferometer (b). Each of the two square pulses exchange population between levels $|11\rangle$ and $|02\rangle$ and thereby can be thought of as beam-splitters whose transmission and reflection coefficients can be tuned by the parameters Δ and $t_p/2$. While waiting for t_{mid} at the sweetspot, the states of the two qubits accrue a relative phase φ . Ideal gate operation achieves full transmission at each beam-splitter such that an arbitrary conditional phase can be implemented by tuning t_{mid} . (c) Landscapes of leakage, L_1 , and conditional phase as a function of Δ and t_{mid} obtained from analytical solving for the dynamics of the gate (with $t_p = \pi/J_2$). (d) Conditional oscillation experiment used to measure the phase of the high-frequency qubit Q_H for the two states of the low-frequency qubit Q_L . (e) Measured data for the experiment for a calibrated gate. Leakage of the gate can be estimated from the average of the two curves of state population of the low frequency qubit using $2L_1 = P_{\text{exc}}^{|1\rangle} - P_{\text{exc}}^{|0\rangle}$. (f) Measured landscapes of (c) using conditional oscillation experiment.

where $\Omega = \sqrt{4J_2^2 + \tilde{\Delta}^2}$ and the parameter ϕ_{dist} accounts for any remaining flux-pulse distortion. The key parameters of interest, namely Δ_0 and,

$$\frac{t_p}{2} = \frac{\pi + \phi_{\text{dist}}}{2J_2}, \quad (2.3)$$

will be crucial in the following subsections. Fitting the landscapes (Fig. Figure 2.5c and Figure 2.5d) allows us to accurately extract the parameters of interest by coarsely sampling over t and Δ_H . Notably, this approach significantly reduces the time required for data acquisition compared to alternative strategies relying on fine sampling.

2.3.3 Sudden net-zero gate

Now, I introduce the sudden net-zero [11] (SNZ) parametrization, as depicted in Fig. Figure 2.6a, for the implementation of two-qubit gates. This parametrization consists of two identical square pulses, each with a duration of $t_p/2$ and equal, but opposite, amplitudes of detuning Δ , separated by an interval of duration t_{mid} . The bipolar nature of this approach

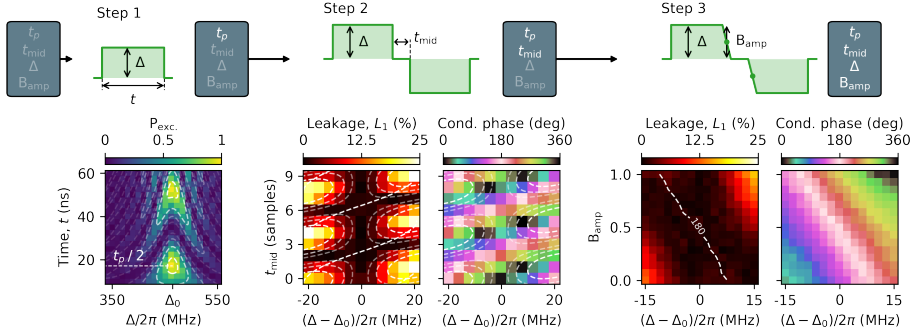


Figure 2.7: Two-qubit gate calibration procedure. Calibration of the four gate parameters: Δ , t_p , t_{mid} , and B_{amp} . Initially, a Chevron experiment determines the best values of t_p and Δ . In the subsequent step, a sweep of t_{mid} and Δ is performed, while simultaneously measuring leakage (L_1) and conditional phase (ϕ_C). The objective here is to identify a suitable value for t_{mid} . In the final step, a similar measurement sweeping Δ and B_{amp} is performed to find the most effective configuration of these parameters. Here, we seek minimize leakage while achieving a 180-degree conditional phase.

mitigates dephasing arising from low-frequency flux noise and also ensures resilience to long-term pulse distortions [9, 28]. However, it's important to note that the SNZ parametrization is most effective when applied to transmons operating at the flux-sweetspot. The operation of the gate relies on interference between states $|11\rangle$ and $|02\rangle$ to realize an arbitrary conditional phase gate. An analogy to an interferometer, as illustrated in Fig. Figure 2.6b, helps illuminate the mechanics of the gate. The parameters t_p and Δ determine the transmission and reflection through the beam splitters, while t_{mid} sets the relative phase accrued, denoted as φ , between paths. Each of these three steps implements a unitary transformation in this manifold as depicted in Fig. Figure A.1 of the appendix section Section A.1. In the ideal scenario, full transmission at the beam-splitter allows the phase of $|11\rangle$ to be solely determined by the relative phase accrued during t_{mid} . To achieve this, each square pulse must complete full exchange between $|11\rangle$ and $|02\rangle$. This condition is met in equation Equation (2.3), where unitary evolution transforms $|11\rangle$ to $U|11\rangle = -i|02\rangle$. Even for slight deviations from this ideal behavior, low leakage can still be achieved through interference [11]. The landscapes of leakage (L_1) and conditional phase (ϕ_C) as functions of Δ and t_{mid} are shown in Fig. Figure 2.6c for unitary evolution of the gate. We experimentally probe these parameters using the Ramsey type experiment illustrated in Fig. Figure 2.6d [9, 11]. Figure Figure 2.6f shows measured results of L_1 and ϕ_C using this experiment. To find optimal parameters for each gate, we use the calibration procedure outlined in Figure Figure 2.7. The calibration process unfolds in three steps. In the first step, we employ the method described in the previous section to determine the pulse duration ($t_p/2$) that maximizes the transfer of $|11\rangle$ to $|02\rangle$.

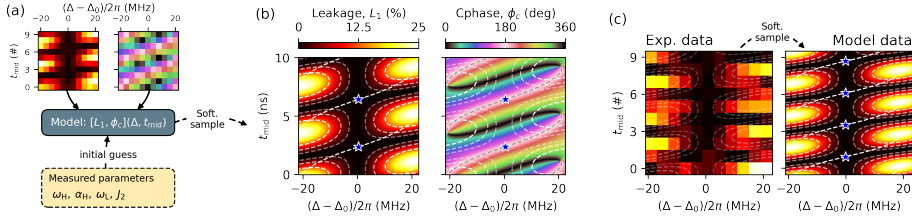


Figure 2.8: SNZ landscape fitting procedure. (a) Measured landscapes of leakage, L_1 , and conditional phase, ϕ_C , are used to fit the model in section Section A.1. Measured parameters are given to the initial guess of the fit to ensure convergence of the high number of free parameters. (b) The same landscapes are sampled in software in order to get much finer detail on the landscape and thus find suitable values of t_{mid} (blue stars). (c) Example of the procedure performed on a landscape where the vertical period of leakage is comparable to the sampling rate of t_{mid} .

Having found t_p , we replicate the landscapes of Fig. Figure 2.6a. Due to the periodic structure of these landscapes, there can be multiple suitable values of t_{mid} . This periodicity is determined by the rate of relative phase accrual (φ in Fig. Figure 2.6b), which, in turn, is dictated by the detuning between levels $|11\rangle$ and $|02\rangle$ at the sweetspot. Since the sweep of t_{mid} is constrained to the AWG sampling rate, we introduce an additional parameter, B_{amp} , to the pulse. This parameter tunes the relative amplitude of a single AWG sampling point between the pulses (Fig. Figure 2.7), thereby finely adjusting the rate of relative phase accrual. Using this method, we first find a value of t_{mid} lying at the intersection between the conditional phase contour of 180 and the leakage fringe at Δ_0 . Once this is done, we perform a fine sweep of Δ and B_{amp} to find a configuration that maximizes gate performance (low L_1 and $\phi_C = 180$). Although conceptually straightforward, achieving the first step in an automated manner can be challenging. Particularly due to the complex features of the landscape and its variable periodicity in the parameter space. To execute this step reliably and efficiently, we employ an analytical model of the landscape and fit it to the measured data. The details of the model are described in the appendix section Section A.1. Similarly to the previous step, this allows extracting the parameters of interest by coarsely sampling over the landscape. We then find all suitable values of t_{mid} by finely sampling from the model in software as shown in Fig. Figure 2.8b. From these, we choose the lowest suitable value. Due to the elevated number of parameters in the model, we ensure reliable convergence of fit by providing an accurate initial guess. This becomes important when analysing landscapes containing many vertical periods of leakage as shown in Figure Figure 2.8c. In these cases, choosing a suitable t_{mid} with the available data might pose a challenge. However, using the model-based interpolation allows performing this step reliably. The runtime for the calibration procedure described in Figure Figure 2.7 is about 24 minutes (Fig. Figure 2.9d) per gate. Despite its

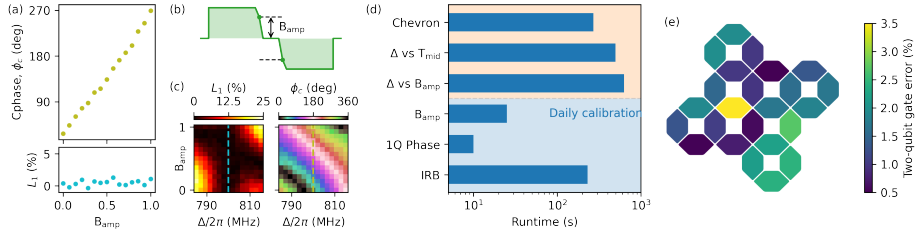


Figure 2.9: **Two-qubit gate daily calibration and performance monitor.** (a) Conditional phase (ϕ_c) and leakage (L_1) as functions of B_{amp} (illustrated in b) along the vertical contour of low leakage in the sudden net-zero landscapes (shown in landscapes c). (d) Runtime of the two-qubit gate calibration routines. The orange shaded area represents the routines depicted in Figure Figure 2.7, while the blue shaded area represents the routines used for daily calibration. These include the B_{amp} sweep (a), single-qubit phase calibration, and interleaved randomized benchmarking.

reliability, its long duration makes it only suitable for initial calibration and not for daily tune-up. Daily calibration only requires correcting for small parameter drift. The runtime for such procedure would desirably be only a few minutes so that we can run through all 24 gates in our device in a couple of hours. For this purpose, we can make use of the vertical valley of leakage of the landscape (Fig. Figure 2.9c) and perform instead a single sweep of B_{amp} as shown in Figure Figure 2.9a. This technique takes ≈ 30 seconds and is sufficient to correct for small drifts in time. Together with single-qubit phase calibration and interleaved randomized benchmarking [27], this procedure has a combined runtime of ≈ 4.5 minutes allowing for full calibration and benchmark of the 24 two-qubit gates (Fig. Figure 2.9e) in under two hours.

2.4 Dealing with interacting two-level systems

In addition to the typical noise sources inherent to the flux-tunable transmon (eg. flux noise), our primary challenge in calibration is unexpected interactions that arise within the tunable bandwidth of the transmon. These are known two-level system (TLS) [29] interactions. The origin of TLSs has been linked to material defects in Josephson junctions [30–33]. Research suggests that TLSs efficiently decay via phonons [34, 35] which makes them generally lossier than typical transmons. Consequently, when TLSs couple to the transmon via its electric dipole, they typically enhance its relaxation, influencing the local $\Gamma_{T_1}(\omega)$ spectrum. In more challenging scenarios, TLSs can exhibit even strong coupling with the transmon, leading to coherent exchange between the two systems. Such examples can be seen in Figure Figure 2.10. While efforts have been made to reduce the density of TLSs using design [31, 32], material selection [36], and material surface treatment [37], current transmon processors

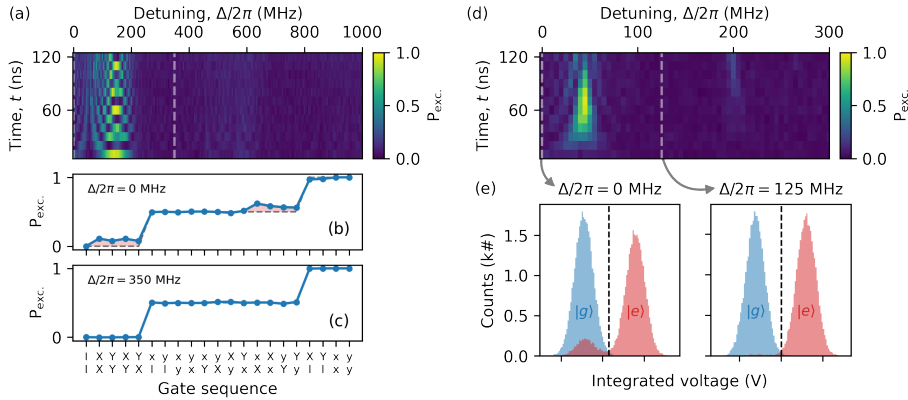


Figure 2.10: Impact of TLS interactions on single-qubit gates and readout. (a) When a strongly coupled TLS interacts with the transmon at its sweet-spot ($\Delta/2\pi = 0$ MHz), coherent control of the qubit using single-qubit gates is hindered, as demonstrated by the allXY sequence (b). Detuning the qubit to a lower frequency ($\Delta/2\pi = 350$ MHz, denoted by the dashed curve) allows for coherent qubit control (c). (d) In this case, the TLS is farther away from resonance with the transmon at the sweet-spot. However, relaxation effects are evident in the single-shot readout histograms (e). This phenomenon can be explained by the AC-Stark shift induced by the readout, which brings the transmon closer to resonance with the TLS. Detuning the transmon beyond this interaction point ($\Delta/2\pi = 150$ MHz) resolves this issue.

have to contend with TLSs [38, 39]. Due to the random nature of these defects, they can impact the operation of qubits in various ways. This impact will depend on their resonance frequency relative to the qubits. In this section, we outline different resonance scenarios of TLSs and techniques to mitigate this impact on qubit operations.

2.4.1 Impact of TLSs in single-qubit gates and readout

When a TLS is resonant with the transmon at its idle frequency, (usually its flux sweet-spot), it disrupts the operation of single-qubit gates (Figs. Figure 2.10a- Figure 2.10c). If resonance occurs just below this frequency, single-qubit gates might not be affected. However, the AC-Stark shift induced by the readout can lead to significant relaxation instead [40] (Figs. Figure 2.10d- Figure 2.10f). To mitigate this, one solution is to operate the transmon at lower frequencies outside its sweet-spot. However, this comes at the cost of increased flux noise and, consequently, lower coherence time.

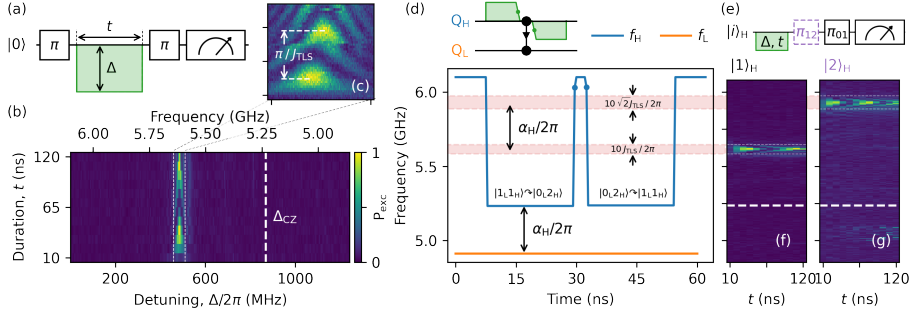


Figure 2.11: Off-resonant TLS in two-qubit gate trajectories. (a) Quantum circuit used for time-domain characterization of a TLS interaction. (b) Excited state probability P_{exc} as function of detuning Δ and duration t . In this interaction landscape, we observe a TLS situated between the idle sweet-spot frequency and the two-qubit gate interaction frequency, Δ_{CZ} , of transmon, Q_H , with another transmon, Q_L , with lower transition frequency. (c) Detailed characterization of the TLS interaction chevron with a coupling strength of J_{TLS} . It is evident that the TLS exhibits frequency instability throughout the measurement. (d) Frequency trajectory of Q_H and Q_L during the two-qubit gate and interaction landscapes of Q_H (f, g). This trajectory highlights the resonance of the TLS concerning the $|0\rangle \rightarrow |1\rangle$ (f) and $|1\rangle \rightarrow |2\rangle$ (g) transitions of the transmon. The procedure used for measuring these landscapes is similar to the one described in (a). The red shaded areas in (d) indicate resonance intervals with the TLS, corresponding to $10J_{TLS}$ and $10\sqrt{2}J_{TLS}$ for the $|0\rangle \rightarrow |1\rangle$ and $|1\rangle \rightarrow |2\rangle$ transitions, respectively.

2.4.2 Impact of TLSs in two-qubit gates

In other cases, such interactions may not overlap with the transmon at its idle frequency. However, they must be crossed during the frequency excursion for two-qubit gates, as illustrated in Figure Figure 2.11. When this situation arises, diabatic transitions can occur between the two systems. The diabatic population lost to the TLS, P_{LZ} , is given by the Landau-Zener formula,

$$P_{LZ} = e^{-2\pi\Gamma}, \Gamma = \frac{J_{TLS}^2}{\hbar \left| \frac{d\Delta}{dt} \right|}, \quad (2.4)$$

where J_{TLS} is the coupling between the states crossing of the transmon and the TLSs. To minimize population loss in the transmon, a possible strategy might involve relying on coherent exchange between the two systems. However, as shown in Figure Figure 2.11c, we observe that such interactions are not stable over time, making them unreliable. Instead, our

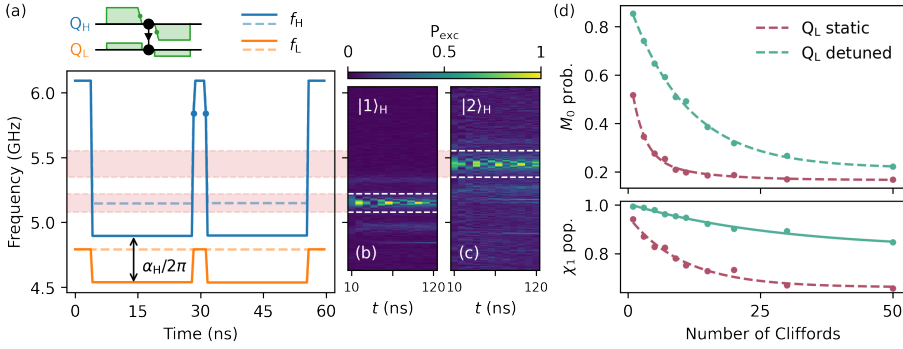


Figure 2.12: Avoiding TLS resonant with two-qubit gate interaction. (a) Frequency trajectories of a high-frequency transmon, Q_H , and a low-frequency transmon, Q_L , respectively, during a two-qubit gate. (b, c) TLS interaction landscapes of Q_H for initial states $|1\rangle$ and $|2\rangle$, showing a strongly coupled TLS lying at the interaction frequency $\omega_H = \omega_L - \alpha_H$. The dashed curves in (a) denote the trajectory of the two-qubit gate with Q_L static at its idle frequency, while the solid curves denote the two-qubit trajectory where Q_L is dynamically detuned to avoid the TLS. (d) Two-qubit Clifford randomized benchmarking of the gate using each of the two-qubit trajectories. From interleaved randomized benchmarking, one obtains errors of 7% and 3% for Q_L in the static and detuned states, respectively.

primary strategy is to minimize P_{LZ} by crossing the TLS as rapidly as possible, satisfying the condition:

$$\left| \frac{d\Delta}{dt} \right| \gg J_{\text{TLS}}. \quad (2.5)$$

When doing this, it is important to consider all relevant transmon levels that cross the TLSs, as depicted in Figure 2.11d. This consideration becomes particularly important during the sudden-net zero gate, where there is population exchange to the non-computational state $|2\rangle$. Due to the anharmonic nature of the transmon, the same TLSs will be resonant with the $|1\rangle \rightarrow |2\rangle$ transition at a detuning $\Delta - \alpha$. This is observed in Figures 2.11f and Figure 2.11g. Additionally, the bandwidth of this interaction is now determined by $\sqrt{2}J_{\text{TLS}}$ since it occurs between states in the second excitation manifold. Therefore, to ensure that every relevant level crosses the TLS as sudden as possible (Eq. Equation (2.5)), we choose B_{amp} such that it does not overlap with the two resonance conditions (shaded area in Fig. Figure 2.11) or the area between them. In other instances, the TLS directly resonates with the interaction frequency of the two-qubit gate, as depicted in Figure 2.12. In such scenarios, we employ a dynamic approach to change the interaction frequency by detuning the low-frequency qubit throughout the gate. This is performed using a net-zero flux pulse applied to the low-frequency qubit, as demonstrated in Figure 2.12a. The gate calibration for such a procedure follows the same steps described in the previous section. It's important

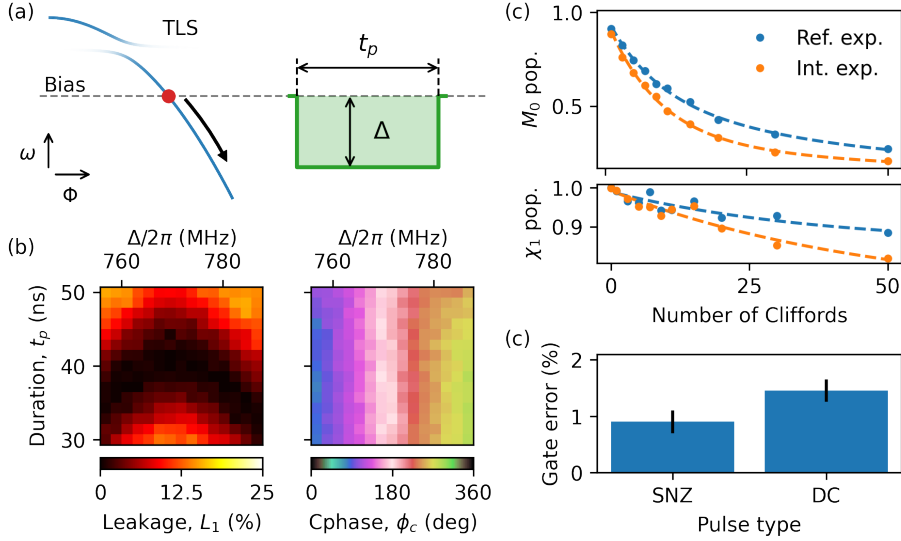


Figure 2.13: **Unipolar two-qubit gates.** (a) Schematic of transmon transition frequency versus flux. To avoid the TLS, shown by the avoided crossing, the transmon is biased to an intermediate frequency (red dot). In this configuration, a unipolar DC pulse (green) is used to perform two-qubit gates. (b) Landscapes of leakage, L_1 , and conditional phase, ϕ_c , measured as function of the DC pulse parameters, Δ and t_p . (c) Interleaved randomized benchmarking for the DC two-qubit gate. The corresponding error and leakage per gate is 1.(5)% and 0.(2)%, respectively. (d) Gate error measured for a sudden net zero (SNZ) and DC two-qubit gates calibrated for the same qubit at the flux sweetspot.

to note that even with this technique, Q_H still crosses the TLS four times, and therefore the same precautions must be taken to ensure Eq. Equation (2.5) is satisfied. Due to the multiple crossings of the TLS and the additional dephasing incurred, gate errors are expected to be significantly higher. However, despite these challenges, we manage to make substantial improvements in gate error rates, achieving moderate error rates as illustrated in Figure Figure 2.12d.

2.4.3 Unipolar two-qubit gates

When operating a transmon at its flux-insensitive sweetspot is not an option due to a nearby TLS, the sudden net-zero (SNZ) pulse parametrization [11] becomes infeasible. To circumvent resonance with the TLS during two-qubit gates, one solution is to use unipolar flux pulses, as depicted in Figure 2.13a. For these applications, a simpler parametrization – the DC pulse – suffices. This pulse is a square pulse with two parameters: amplitude, Δ , and pulse duration,

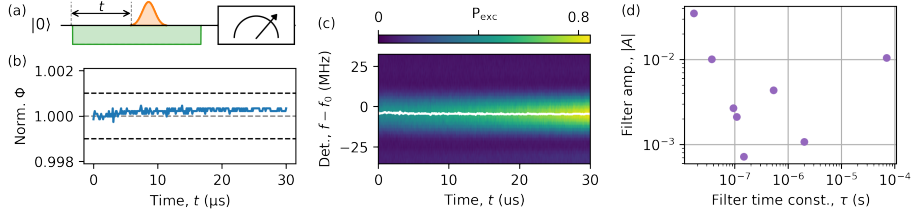


Figure 2.14: Correction of long time-scale flux distortions. (a) Experimental procedure to measure long time-scale flux step response. Here, the step response of a square flux pulse (green) is measured by sweeping the delay, t , and frequency, f , of a microwave pulse (orange) and measuring the qubit population, P_{exc} , after the pulse. (b) Inferred step response from the measured frequency obtained from the resulting P_{exc} landscape (c). The frequency response (white curve) is obtained by fitting P_{exc} for each delay, t . (d) Infinite impulse response (IIR) filter parameters used to obtain the calibrate step response in (b). Each point corresponds to a set of parameters describing the exponential pulse correction modeled by Eq. 2.6.

t_p . An advantage of this parametrization is its simpler and faster calibration procedure, which reduces to the first step shown in Figure 2.7. Figure 2.13b shows the measured leakage and conditional phase landscapes for this step, consisting of a 2D sweep of Δ and t_p . Here, the leakage landscape corresponds to the chevron pattern arising from the resonance of $|11\rangle$ and $|02\rangle$. The contour of $\phi_c = 180$ degrees aligns with this resonance. In this regime, the gate dynamics can be understood as a 2π rotation in the manifold $\{|11\rangle, |02\rangle\}$, described by the Hamiltonian in Eq. 2.1. This unitary imparts a phase of $e^{i\pi}$ to both $|11\rangle$ and $|02\rangle$, in addition to a single-qubit phase on the high-frequency qubit.

Since this parametrization is no longer net zero, an important development required to operate these gates is correcting flux pulse distortions on longer $\sim \mu\text{s}$ time scales. At these scales, the Ramsey-based cryoscope approach [28] is not viable, as qubit dephasing leads to loss of signal-to-noise ratio (SNR). We address this using the spectroscopy-based measurement shown in Figure 2.14a. This method reconstructs the flux step response (Fig. 2.14b) by measuring the qubit population as a function of the time and frequency of a microwave pulse (Fig. 2.14c). Similar to the cryoscope technique, we fit the normalized step response, $S(t)$, to the model

$$S(t) = g(1 + Ae^{-t/\tau}), \quad (2.6)$$

where the amplitude, A , and characteristic time, τ , parameters are used to parameterize real-time infinite impulse response (IIR) filters. Calibration is performed iteratively, using up to eight filters. Figure 2.14d shows an example of these parameters after calibrating the step response shown in Fig. 2.14b. Having calibrated for long time scale flux distortions, we perform interleaved randomized benchmarking of the DC two-qubit gate (Fig. 2.13d). We find a 1.(6)% error rate which is comparable to other gates in the device. To compare the per-

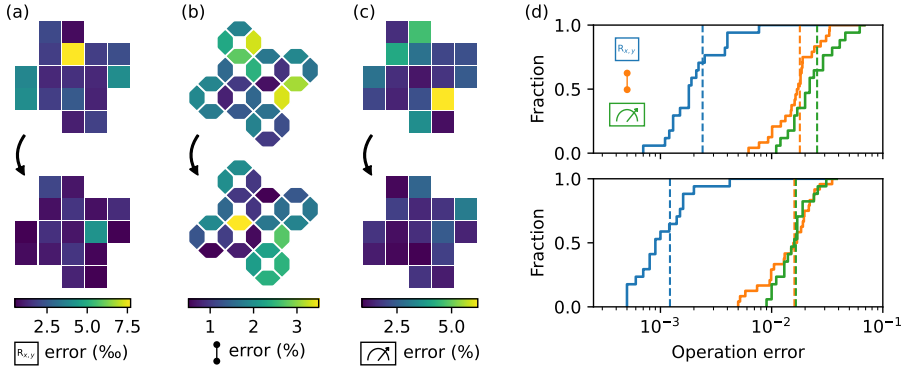


Figure 2.15: **Comparison of error rates for different devices.** Different error rates evaluated for two different Surface-17 devices of different versions. (a) Single-qubit gate error evaluated using Clifford randomized benchmarking. (b) Two-qubit gate error evaluated using Interleaved randomized benchmarking. (c) Single-shot readout assignment fidelity. (d) Cumulative histogram of the aforementioned error rates for each device. The mean error rate is denoted by the dashed vertical line.

performances of SNZ and DC pulse parametrizations, we benchmarked two-qubit gates using each parametrization for the same qubit biased at the sweetspot (Fig. 2.13d). As expected, the gate error is higher for the DC pulse (0.(9)% and 1.(5)% for SNZ and DC, respectively). To summarize, unipolar pulses enable qubit operation even when a TLS is resonant at the sweetspot. However, this approach incurs increased dephasing during the gate because the pulse is no longer symmetric with respect to the sweetspot, making the transmon more susceptible to low-frequency flux noise when compared to the SNZ gate. Consequently, this configuration is used only as a last resort when operating the transmon at the sweetspot is overly detrimental.

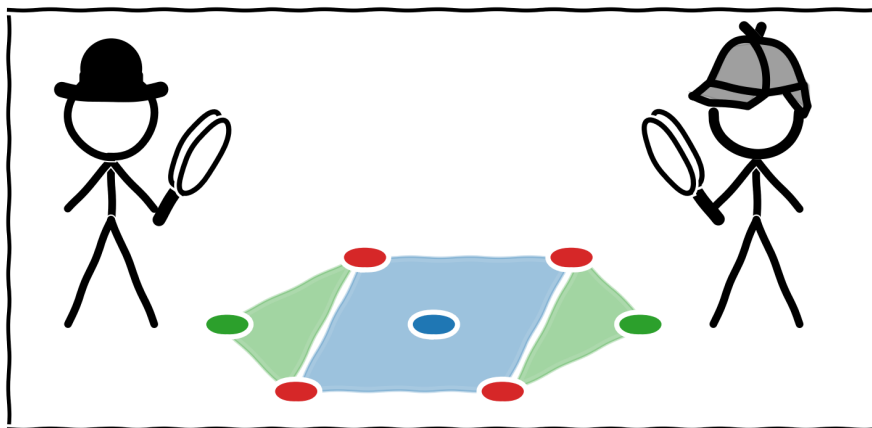
2.5 Assessing performance and progress

Metrics like gate fidelity are indicative of the overall performance of the device. Although they do not account for crosstalk errors, they are still useful during the development process. This is particularly relevant when iterating through different devices after improvements in, for example, the fabrication recipe, and we want to assess how these improvements translate into processor performance during an algorithm. Figure 2.15 shows single- and two-qubit gate errors along with readout assignment errors for two Surface-17 devices. Each device represents a different iteration where improvements were made in both fabrication and parameter design. These improvements notably resulted in higher relaxation times, T_1 (increasing from 10 to $27\mu s$ on average across all qubits), and better targeting of resonator parameters [41].

From the cumulative distribution of errors shown in Figure 2.15d, the most substantial reductions are observed in single-qubit gate and readout assignment errors. We attribute the former to the increased relaxation times and the latter to the more precise targeting of parameters. The average two-qubit gate error shows only a marginal improvement. We speculate that this is because these gates are primarily limited by the TLS error mechanisms discussed in the previous section.

LOGICAL QUBIT OPERATIONS IN AN ERROR DETECTING SURFACE CODE

J. F. Marques, B. M. Varbanov, M. S. Moreira, H. Ali, N. Muthusubramanian, C. Zachariadis, F. Battistel, M. Beekman, N. Haider, W. Vlothuizen, A. Bruno, B. M. Terhal & L. DiCarlo



Parts of this chapter have been published in *Nat. Phys.* **18**, 80-86 (2022).

3.1 Introduction

TWO key capabilities will distinguish an error-corrected quantum computer from present-day noisy intermediate-scale quantum (NISQ) processors [42]. First, it will initialize, transform and measure quantum information encoded in logical qubits rather than physical qubits, where a logical qubit is a highly entangled two-dimensional subspace in the larger Hilbert space of many more physical qubits. Second, it will use repetitive quantum parity checks to discretize, signal and (with the aid of a decoder) correct errors occurring in the constituent physical qubits without destroying the encoded information [43]. Provided the incidence of physical errors is below a code-specific threshold and the quantum circuits for logical operations and stabilization are fault tolerant, the logical error rate can be exponentially suppressed by increasing the distance (redundancy) of the quantum error correction (QEC) code employed [44]. The exponential suppression for specific physical qubit errors (bit-flip or phase-flip) has been experimentally demonstrated [45–48] for repetition codes [49–51].

Leading experimental quantum platforms have taken key steps towards implementing QEC codes protecting logical qubits from general physical qubit errors. In particular, trapped-ion systems have demonstrated logical-level initialization, gates and measurements for single logical qubits in the Calderbank-Shor-Steane [52] and Bacon-Shor [53] codes. Most recently, entangling operations between two logical qubits have been demonstrated in the surface code using lattice surgery [54]. However, except for smaller-scale experiments using two ion species [55], trapped-ion experiments in QEC have so far been limited to a single round of stabilization.

In parallel, taking advantage of highly-non-demolition measurement in circuit quantum electrodynamics [56], superconducting circuits have taken key strides in repetitive stabilization of two-qubit entanglement [57, 58] and logical qubits. Quantum memories based on 3D-cavity logical qubits in cat [59, 60] and Gottesman-Kitaev-Preskill [61] codes have crossed the memory break-even point. Meanwhile, monolithic architectures have focused on logical qubit stabilization in a surface code realized with a 2D lattice of transmon qubits. Currently, the surface code [62] is the most attractive QEC code for solid-state implementation owing to its practical nearest-neighbor-only connectivity requirement and high error threshold. Recent experiments [46, 63] have demonstrated repetitive stabilization by post-selection in a surface code which, owing to its small size, is capable of quantum error detection but not correction. In particular, Ref. [63] has demonstrated the preparation of logical cardinal states and logical measurement in two cardinal bases. Here, we go beyond previous work by demonstrating a complete suite of logical-qubit operations for this small (distance-2) surface code while preserving multi-round stabilization. Our logical operations include initialization anywhere on the logical Bloch sphere (with significant improvement over previously reported fidelities), measurement in all cardinal bases, and a universal set of single-qubit logical gates. For each type of operation, we quantify the increased performance of fault-tolerant variants over non-fault-tolerant ones. We use a logical Pauli transfer matrix to describe a logical gate, analogous to

the procedure commonly used to describe gates on physical qubits [64]. Finally, we perform logical state stabilization by means of repeated error detection where we compare the performance of two scalable, fault-tolerant stabilizer measurement schemes compatible with our quantum hardware architecture [21].

The distance-2 surface code (Fig. Figure 3.1a) uses four data qubits (D_1 through D_4) to encode one logical qubit, whose two-dimensional codespace is the even-parity (i.e., eigenvalue +1) subspace of the stabilizer set

$$\mathcal{S} = \{Z_{D1}Z_{D3}, X_{D1}X_{D2}X_{D3}X_{D4}, Z_{D2}Z_{D4}\}. \quad (3.1)$$

This codespace has logical Pauli operators

$$Z_L = Z_{D1}Z_{D2}, Z_{D3}Z_{D4}, Z_{D1}Z_{D4}, \text{ and } Z_{D2}Z_{D3}, \quad (3.2)$$

$$X_L = X_{D1}X_{D3} \text{ and } X_{D2}X_{D4}, \quad (3.3)$$

that anti-commute with each other and commute with $\mathcal{S}$, and logical computational basis

$$|0_L\rangle = \frac{1}{\sqrt{2}} (|0000\rangle + |1111\rangle), \quad (3.4)$$

$$|1_L\rangle = \frac{1}{\sqrt{2}} (|0101\rangle + |1010\rangle). \quad (3.5)$$

Measuring the stabilizers using three ancilla qubits (A_1, A_2 and A_3 in Fig. Figure 3.1a) allows detection of all physical errors that change the outcome of one or more stabilizers to $m = -1$. This list includes all errors on any one single qubit. However, no error syndrome is unique to a specific physical error. For instance, a phase flip in any one data qubit triggers the same syndrome: $m_{A2} = -1$. Consequently, this code cannot be used to correct such errors. We thus perform state stabilization by post-selecting runs in which no error is detected by the stabilizer measurements in any cycle. In this error-detection context, an operation is fault-tolerant if any single-fault produces a non-trivial syndrome and can therefore be post-selected out [65] (see Suppl. Material).

3.2 Results

3.2.1 Stabilizer measurements

Achieving high performance in a code hinges on performing projective quantum parity (stabilizer) measurements with high assignment fidelity, meaning one can accurately discriminate parity, and low additional backaction such that the state of the qubits after the measurement is properly projected onto the parity subspace. We implement each of the stabilizers in $\mathcal{S}$ using a standard indirect-measurement scheme [66, 67] with a dedicated ancilla. We benchmark the accuracy of each parity measurement by preparing the data-qubits in a computational state and measuring the probability of ancilla outcome $m_A = -1$. As a fidelity metric, we calculate the average probability to correctly assign the parity $Z_{D1}Z_{D3}$, $Z_{D1}Z_{D2}Z_{D3}Z_{D4}$ and $Z_{D1}Z_{D3}$, finding 94.2%, 86.1% and 97.2%, respectively (see Suppl. Material Fig. S2).

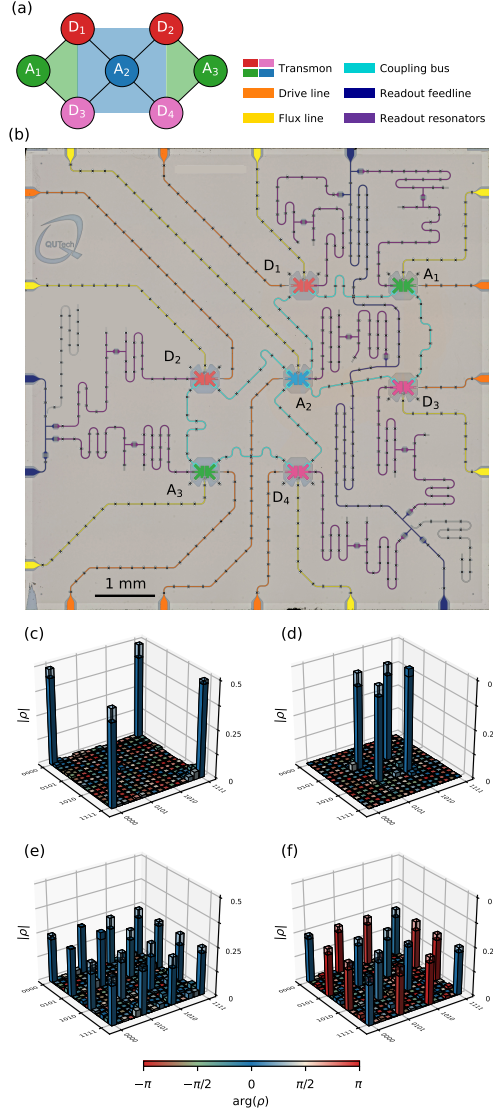


Figure 3.1: **Surface-7 quantum processor and initialization of logical cardinal states.**

(a) Distance-two surface code. (b) Optical image of the quantum hardware with added false-color to emphasize different circuit elements. (c-f) Estimated physical density matrices, ρ , after targeting the preparation of the logical cardinal states $|0_L\rangle$ (c), $|1_L\rangle$ (d), $|+_L\rangle$ (e) and $|-_L\rangle$ (f). Each state is measured after preparing the data qubits in $|0000\rangle$, $|1010\rangle$, $|++++\rangle$ and $|++--\rangle$, respectively. The ideal target state density matrix is shown in the shaded wireframe.

3.2.2 Logical state initialization using stabilizer measurements

A practical means to quantify the backaction of stabilizer measurements is using them to initialize logical states. As proposed in Ref. [63], we can prepare arbitrary logical states by first initializing the data-qubit register in the product state

$$|\psi\rangle = \left(C_{\theta/2}|0\rangle + S_{\theta/2}|1\rangle\right)|0\rangle \left(C_{\theta/2}|0\rangle + S_{\theta/2}e^{i\phi}|1\rangle\right)|0\rangle \quad (3.6)$$

using single-qubit rotations R_y^θ on D_1 and R_ϕ^θ on D_3 acting on $|0000\rangle$ ($C_\alpha = \cos \alpha$ and $S_\alpha = \sin \alpha$). A follow-up round of stabilizer measurements ideally projects the four-qubit state onto the logical state

$$|\psi_L\rangle = \left(C_{\theta/2}^2|0_L\rangle + S_{\theta/2}^2e^{i\phi}|1_L\rangle\right) / \sqrt{C_{\theta/2}^4 + S_{\theta/2}^4} \quad (3.7)$$

with probability

$$P = \frac{1}{2} \left(C_{\theta/2}^4 + S_{\theta/2}^4\right). \quad (3.8)$$

We use this procedure to target initialization of the logical cardinal states $|0_L\rangle$, $|1_L\rangle$, $|+_L\rangle = (|0_L\rangle + |1_L\rangle)/\sqrt{2}$, and $|-_L\rangle = (|0_L\rangle - |1_L\rangle)/\sqrt{2}$. For the first two states, the procedure is fault-tolerant according to the definition above. We characterize the produced states using full four-qubit state tomography including readout calibration and maximum-likelihood estimation (MLE) (Fig. Figure 3.1c-f). The fidelity F_{4Q} to the ideal four-qubit target states is $90.0 \pm 0.3\%$, $92.9 \pm 0.2\%$, $77.3 \pm 0.5\%$, and $77.1 \pm 0.5\%$, respectively. For each state, we can extract a logical fidelity F_L by further projecting the obtained four-qubit density matrix onto the codespace [63], finding $99.83 \pm 0.08\%$, $99.97 \pm 0.04\%$, $96.82 \pm 0.55\%$, and $95.54 \pm 0.55\%$, respectively (see Methods). This sharp increase from F_{4Q} to F_L demonstrates that the vast majority of errors introduced by the parity check are weight-1 and detectable. A simple modification makes the initialization of $|+_L\rangle$ ($|-_L\rangle$) also fault-tolerant: initialize the data-qubit register in a different product state, namely $|++++\rangle$ ($|++--\rangle$), before performing the stabilizer measurements. With this modification, F_{4Q} increases to $85.4 \pm 0.3\%$ ($84.6 \pm 0.3\%$) and F_L to $99.78 \pm 0.09\%$ ($99.64 \pm 0.17\%$), matching the performance achieved when targeting $|0_L\rangle$ and $|1_L\rangle$.

3.2.3 Logical measurement of arbitrary states

A key feature of a code is the ability to measure logical operators. In the surface code, we can measure X_L (Z_L) fault-tolerantly, albeit destructively, by simultaneously measuring all data qubits in the X (Z) basis to obtain a string of data-qubit outcomes (each $+1$ or -1). The value assigned to the logical operator is the computed product of data-qubit outcomes as prescribed by Eq. Equation (3.3) (Equation (3.2)). Additionally, the outcome string is used to compute a value for the stabilizer(s) $X_{D1}X_{D2}X_{D3}X_{D4}$ ($Z_{D1}Z_{D3}$ and $Z_{D2}Z_{D4}$), enabling a final step of error detection (Fig. Figure 3.2a). Measurement of $Y_L = +iX_LZ_L = Y_{D1}Z_{D2}X_{D3}$

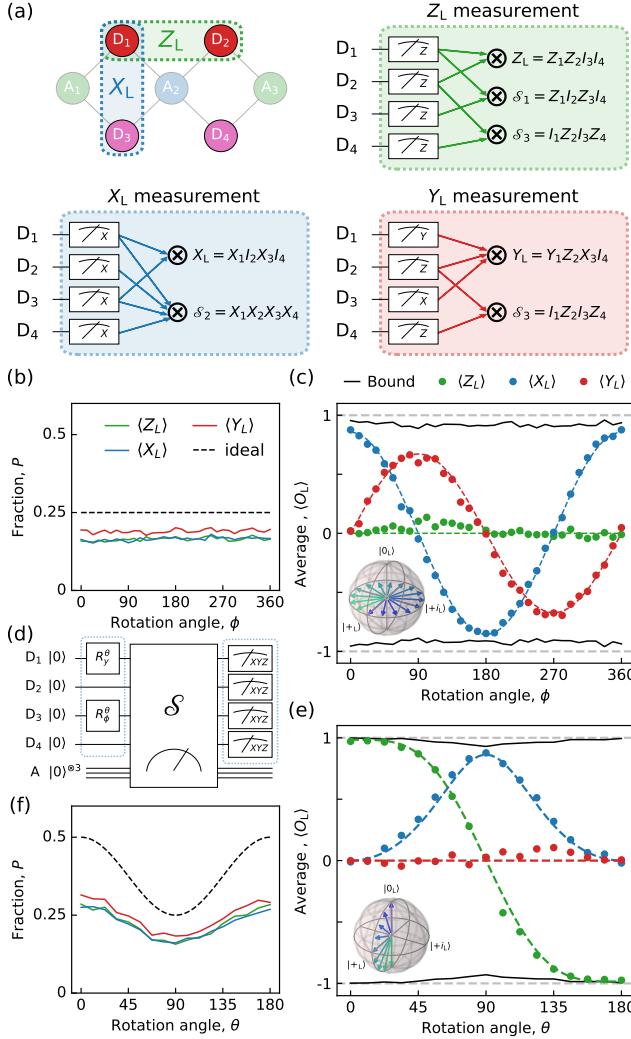


Figure 3.2: **Arbitrary logical-state initialization and measurement in the logical cardinal bases.** (a) Assembly of data-qubit measurements used to evaluate logical operators Z_L , X_L and Y_L with additional error detection. (d) Initialization of logical states using the procedure described in Eq. 3.6. (c, e) Z_L , X_L and Y_L logical measurement results as a function of the gate angles ϕ (c) and θ (e). The colored dashed curves show a fit of the analytical prediction based on Eqs. 3.9 and 3.11 to the data and the dark curve denotes a bound based on the measured F_L of each state. (b, f) Total fraction P of post-selected data as a function of the input angle for each logical measurement. The dashed curve shows the ideal fraction given by Eq. 3.8.

is not fault-tolerant. However, we lower the logical assignment error by also measuring D_4 in the Z basis to compute a value for $Z_{D_2}Z_{D_4}$ and thereby detect bit-flip errors in D_2 and D_4 .

We demonstrate Z_L , X_L and Y_L measurements on logical states prepared on two orthogonal planes of the logical Bloch sphere. Setting $\theta = \pi/2$ and sweeping ϕ , we ideally prepare logical states on the equator (Fig. Figure 3.2d)

$$|\psi_L\rangle = (|0_L\rangle + e^{i\phi}|1_L\rangle) / \sqrt{2}. \quad (3.9)$$

We measure the produced states in the Z_L , X_L and Y_L bases and obtain experimental averages $\langle Z_L \rangle$, $\langle X_L \rangle$ and $\langle Y_L \rangle$. As expected, we observe sinusoidal oscillations in $\langle X_L \rangle$ and $\langle Y_L \rangle$ and near-zero $\langle Z_L \rangle$. The reduced range of the $\langle Y_L \rangle$ oscillation evidences the non-fault-tolerant nature of Y_L measurement. A second manifestation is the higher fraction P of post-selected data for Y_L (Fig. Figure 3.2b). To quantify the logical assignment fidelity F_L^R with correction for initialization error, it is tempting to apply the formula

$$\frac{\langle O_L \rangle_{\max} - \langle O_L \rangle_{\min}}{2} = (2F_L^R - 1)(2F_L - 1), \quad O \in \{X, Y\} \quad (3.10)$$

inspired by the standard method to quantify readout fidelity of physical qubits from Rabi oscillations with limited initialization fidelity (described in the Supplementary Material). This method suggests $F_L^R = 95.8\%$ for X_L and 87.5% for Y_L . However, this method is not accurate for a logical qubit because not all input states outside the codespace are rejected by the limited set of stabilizer checks computable from the data-qubit outcome string and, moreover, detectable initialization errors can become undetectable when compounded with data-qubit readout errors. An accurate method to extract F_L^R based on the measured 16×16 data-qubit assignment probability matrix (detailed in the Supplementary Material) gives $F_L^R = 98.7\%$ for X_L and 91.4% for Y_L .

Setting $\phi = 0$ and sweeping θ , we then prepare logical states on the X_L - Z_L plane (Fig. Figure 3.2e), ideally

$$|\psi_L\rangle = (C_{\theta/2}^2 |0_L\rangle + S_{\theta/2}^2 |1_L\rangle) / \sqrt{C_{\theta/2}^4 + S_{\theta/2}^4}. \quad (3.11)$$

Note that due to the changing overlap of the initial product state with the codespace, P is now a function of θ (Eq. Equation (3.8)). The approximate extraction method based on the range of $\langle Z_L \rangle$ suggests $F_L^R = 99.4\%$, while the accurate method gives 99.8% . Note that while both are fault-tolerant, the Z_L measurement has higher fidelity than the X_L measurement as the former is only vulnerable to vertical double bit-flip errors while the latter is vulnerable to both horizontal and diagonal double phase-flip errors.

3.2.4 Logical gates

Finally, we demonstrate a suite of gates enabling universal logical-qubit control (Fig. Figure 3.3). Full control of the logical qubit requires a gate set comprising Clifford and non-Clifford logical gates. Some Clifford gates, like $R_{Z_L}^\pi$ and $R_{X_L}^\pi$ (where $R_{O_L}^\theta = e^{-i\theta O_L/2}$),

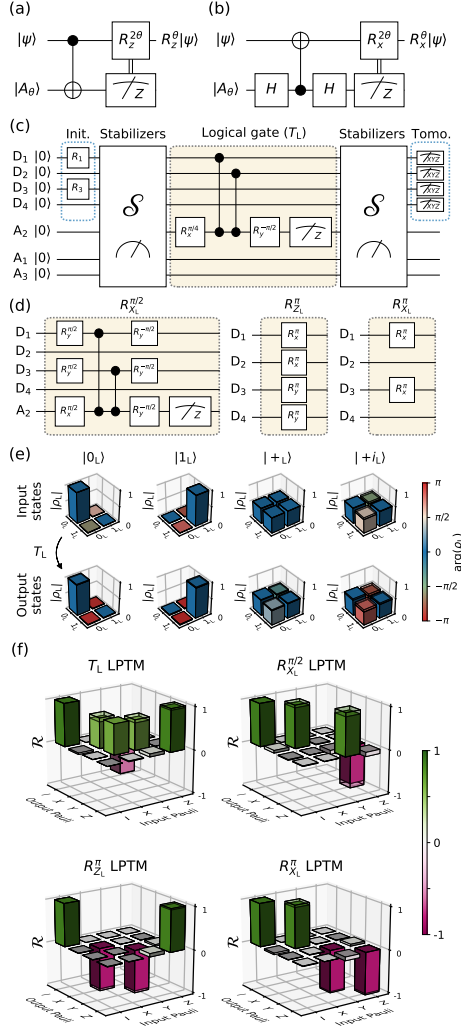


Figure 3.3: Logical gates and their characterization. (a, b) General gate-by-measurement schemes realizing arbitrary rotations around the Z (a) and X (b) axis of the Bloch sphere. (c) Process tomography experiment of the T_L gate. Input cardinal logical states are initialized using the method of Fig. 3.2. Output states are measured following a second round of stabilizer measurements. (d) Logical $R_{X_L}^{\pi/2}$, $R_{Z_L}^{\pi}$ and $R_{X_L}^{\pi}$ gates compiled using our hardware-native gateset. (e) Logical state tomography of input and output states of the T_L gate. These logical density matrices are obtained by performing four-qubit tomography of the data qubits and then projecting onto the codespace. (f) Extracted (solid) and ideal (wireframe) logical Pauli transfer matrices.

can be implemented transversally and therefore fault-tolerantly (Fig. Figure 3.3d). We perform arbitrary rotations (generally non-fault-tolerant) about the Z_L axis using the standard gate-by-measurement circuit [68] shown in Fig. Figure 3.3a. In our case, the ancilla is physical (A_2), while the qubit transformed is our logical qubit. The rotation angle θ is set by the initial ancilla state $|A_\theta\rangle = (|0\rangle + e^{i\theta}|1\rangle)/\sqrt{2}$. Since we cannot do binary-controlled Z_L rotations, we simply post-select runs in which the measurement outcome is $m_{A_2} = +1$. However, we note that these gates can be performed deterministically using repeat-until-success [69]. Choosing $\theta = \pi/4$ implements the non-Clifford $T_L = R_{Z_L}^{\pi/4}$ gate. A similar circuit (Fig. Figure 3.3b) can be used to perform arbitrary rotations around the X_L axis. We compile both circuits using our hardware-native gateset (Figs. Figure 3.3c,d). To assess logical-gate performance, we perform logical process tomography using the procedure illustrated in Fig. Figure 3.3e for T_L . First, we initialize into each of the six logical cardinal states $\{|0_L\rangle, |1_L\rangle, |+_L\rangle, |-_L\rangle, |+i_L\rangle, |-i_L\rangle\}$. We characterize each actual input state by four-qubit state tomography and project to the codespace to obtain a logical density matrix. Next, we similarly characterize each output state produced by the logical gate and a second round of stabilizer measurements to detect errors occurred in the gate (full data in Fig. S3). Using this over-complete set of input-output logical-state pairs, combined with MLE (see Methods), we extract a logical Pauli transfer matrix (LPTM). The resulting LPTMs for the non-fault-tolerant T_L and $R_{X_L}^{\pi/2}$ gates as well as the fault-tolerant $R_{Z_L}^\pi$ and $R_{X_L}^\pi$ are shown in Fig. Figure 3.3e. From the LPTMs, we extract average logical gate fidelities F_L^G (Eq. Equation (3.19)) 97.3%, 95.6%, 97.9%, and 98.1%, respectively.

3.2.5 Pipelined versus parallel stabilizer measurements

A scalable control scheme is fundamental to realize surface codes with large code distance. To this end, we now compare the performance of two schemes suitable for the quantum hardware architecture proposed in Ref. [21]. These schemes are scalable in the sense that their cycle duration remains independent of code distance. The pipelined scheme interleaves the coherent operations and ancilla readout steps associated with stabilizer measurements of type X and Z by performing the coherent operations of X (Z) type stabilizers during the readout of Z (X) type stabilizers (Fig. Figure 3.4a). The parallel scheme performs all ancilla readouts simultaneously (Fig. Figure 3.4b). The pipelined cycle scheme duration is shorter than the parallel scheme by 16% which can potentially increase the performance of the code. This only occurs if the interleaved readout of ancillas does not result in increased measurement-induced dephasing between them. To compare their performance, we initialize and stabilize $|0_L\rangle$ for up to $n = 15$ cycles. We perform refocusing pulses ($R_{\phi_i}^\pi$) on the data qubits to correct for coherent errors during the measurement of ancilla qubits. We also separately calibrate the equatorial rotation axis of this gate for each scheme to extract the best performance. At each n , we take data back-to-back for the two schemes in order to minimize the effect of parameter drift, repeating each experiment up to 256×10^3 times.

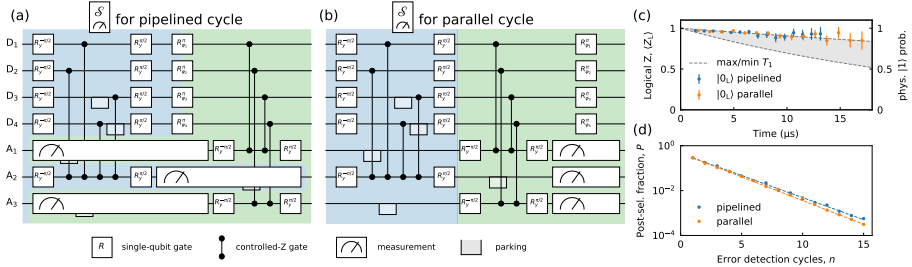


Figure 3.4: Repetitive error detection using pipelined and parallel stabilizer measurement schemes. (a, b) Gate sequences used to implement the pipelined (a) and parallel (b) stabilizer measurement schemes. Gate duration is 20 ns for single-qubit gates, 60 ns for controlled-Z (CZ) gates and parking [21, 57], and 540 ns for ancilla readout. The order of CZs in the $X_{D_1}X_{D_2}X_{D_3}X_{D_4}$ stabilizer (blue shaded region) prevents the propagation of ancilla errors into logical qubit errors [65]. The total cycle duration for the pipelined (parallel) scheme is 840 ns (1000 ns). (c) Estimated Z_L expectation value, $\langle Z_L \rangle$, measured for the $|0_L\rangle$ state versus the duration of the experiment using the pipelined (blue) and the parallel (orange) schemes. We also plot the excited-state probability (right axis) set by the maximum and minimum physical qubit T_1 . (d) Post-selected fraction of data versus the number of error detection cycles n for the pipelined (blue) and parallel (orange) scheme.

Figure 3.4c shows the Z_L measurement outcome averaged over the post-selected runs. We extract the error-detection rate γ from the n -dependence of the fraction of post-selected data P (Fig. 3.4d) using the procedure described in Methods. We observe that the error rate is slightly lower for the pipelined scheme ($\gamma_{\text{pip}} \sim 45\%$), most likely due to the shorter duration of the cycle. This superiority is consistent across different input logical states (see Fig. S4) with an average ratio $\gamma_{\text{pip}}/\gamma_{\text{par}} \sim 97\%$.

3.3 Discussion

We have demonstrated a suite of logical-level initialization, gate and measurement operations in a distance-2 superconducting surface code undergoing repetitive stabilizer measurements. For each type of logical operation, we have quantified the increased performance of fault-tolerant variants over non-fault-tolerant variants. Table 3.1 summarizes all the results. We can initialize the logical qubit to any point on the logical Bloch sphere, with logical fidelity surpassing Ref. [63]. In addition to characterizing initialized states using full four-qubit tomography, we also demonstrate logical measurements in all logical cardinal bases. Finally, we demonstrate a universal single-qubit set of logical gates by performing logical process tomography, using the concept of a logical-level Pauli transfer matrix. As expected, the fidelity of the fault-tolerant gates is higher than the non-fault-tolerant ones. However, one would expect

Logical operation		Characteristic	Logical fidelity metric	value (%)
Init.	$ 0_L\rangle$	FT	F_L	99.83
	$ 1_L\rangle$	FT		99.97
	$ +_L\rangle$	Non-FT/FT		96.82/99.78
	$ -_L\rangle$	Non-FT/FT		95.54/99.64
Meas.	Z_L	FT	F_L^R	99.8
	X_L	FT		98.7
	Y_L	Non-FT		91.4
Gate	$R_{X_L}^\pi$	FT	F_L^G	97.9
	$R_{Z_L}^\pi$	FT		98.1
	$R_{X_L}^{\pi/2}$	Non-FT		95.6
	T_L	Non-FT		97.3

Table 3.1: **Summary of logical initialization, measurement, and gate operations and their performance.** Fault-tolerant operations are labelled FT and non-fault tolerant ones Non-FT. Quoted F_L^R values are those extracted with the accurate method described in the Supplementary Material.

a sharper difference given the typical error rates of the operations involved. We believe this could be due to errors introduced by the stabilizer measurements which might be dominant over the errors of the logical gate itself.

With a view towards implementing higher-distance surface codes using our quantum-hardware architecture [21], we have compared the performance of two scalable stabilization schemes: the pipelined and parallel measurement schemes. In this comparison, two main factors compete. On one hand, the shorter cycle time favors pipelining. On the other, the pipelining introduces extra dephasing on ancilla qubits of one type during readout of the other. The performance of both schemes is comparable, but slightly higher for the pipelined scheme. From detailed density-matrix simulations discussed in the Supplementary Material, we further understand that conventional qubit errors such as energy relaxation, dephasing and readout assignment error alone do not fully account for the net error-detection rate observed in the experiment (see Fig. S10 and also not for the P reduction in Figs. Figure 3.2b,f; see Fig. S11). We believe that the dominant error source is instead leakage to higher transmon states incurred during CZ gates. Our data (Fig. S9) shows that the error detection scheme successfully post-selects leakage errors in both the ancilla and data qubits. Learning to identify these non-qubit errors and to correct them without post-selection is the subject of ongoing research [70–72] and an outstanding challenge in the quest for quantum fault-tolerance with higher-distance superconducting surface codes [73], which to this date have yet to be implemented with repeated error correction.

3.4 Methods

3.4.1 Device

We use a superconducting circuit-QED processor (Fig. Figure 3.1b) featuring the quantum hardware architecture proposed in Ref. [21]. Seven flux-tunable transmons are arranged in three frequency groups: a high-frequency group for D_1 and D_2 ; a middle-frequency group for A_1 , A_2 and A_3 ; and a low-frequency group for D_3 and D_4 . Similar to the device in Ref. [63], each transmon is transversely coupled to its nearest neighbors using a coupling bus resonator dedicated to each pair. This simplest and minimal connectivity minimizes multi-qubit crosstalk. Also, every transmon has a dedicated flux line for two-qubit gating, and a dispersively coupled readout resonator with Purcell filter enabling frequency-multiplexed readout [41, 58] using two feedlines. In contrast to Ref. [63], every transmon has a dedicated microwave drive line for single-qubit gating, avoiding the need to drive any via a feedline and thus reducing driving crosstalk.

All transmons are flux biased to their maximal frequency (i.e., flux sweetspot [74]), where measured qubit relaxation (T_1) and dephasing (T_2^{Echo}) times lie in the range 27–102 ns and 55–117 ns, respectively. Detailed information on the implementation and performance of single- and two-qubit gates for this same device can be found in Ref. [11]. Device characteristics are also summarized in Table S1.

The device was fabricated on a high-resistivity intrinsic Si<100> wafer that was first descummed using UV-ozone cleaner and stripped of native oxides using buffered oxide etch solution (BOE 7:1). The wafer was subjected to vapor of hexamethyldisilazane (HMDS) at 150°C and sputtered with 200 nm of niobium titanium nitride (NbTiN). Post dicing into smaller dies, a layer of hydrogen silsesquioxane (HSQ) was spun and baked at 300°C to serve as an inorganic sacrificial mask for wet etching of NbTiN. This layer was removed post base-patterning steps. The quantum plane was defined using electron-beam (e-beam) lithography of a high-contrast, positive-tone resist spun on top of the NbTiN-HSQ stack. Post development, the exposed region was first dry etched using SF₆/O₂ mixture and then wet etched to remove any residual metal. Dolan-bridge-style Al/AlOx/Al Josephson junctions were then fabricated using standard double-angle e-beam evaporation. Airbridges and crossovers were added using a two-step process. The first step involved patterning galvanic contact using e-beam resist ($\sim 6 \mu\text{m}$ thick) subjected to reflow. In the second step, the airbridges and crossovers were patterned with e-beam evaporated Al (450 nm thick). Finally, the device underwent dicing, resist lift-off and Al wirebonding to a printed circuit board.

3.4.2 State tomography

To perform state tomography on the prepared logical states, we measure the $4^4 - 1$ expectation values of data-qubit Pauli observables, $p_i = \langle \sigma_i \rangle$, $\sigma_i \in \{I, X, Y, Z\}^{\otimes 4}$ (except $I^{\otimes 4}$).

Interleaved with this measurement we also characterize the measurement POVM used to correct for readout errors in p_i . These are then used to construct the density matrix

$$\rho = \sum_{i=0}^{4^4-1} \frac{p_i \sigma_i}{2^4} \quad (3.12)$$

with $p_0 = 1$, corresponding to $\sigma_0 = I^{\otimes 4}$. Due to statistical uncertainty in the measurement, the constructed state, ρ , might lack the physicality characteristic of a density matrix, that is, $\text{Tr}(\rho) = 1$ and $\rho \geq 0$. Specifically, ρ might not satisfy the latter constraint, while the former is automatically satisfied by $p_0 = 1$. To enforce these constraints, we use a maximum-likelihood method [64] to find the physical density matrix, ρ_{ph} , that is closest to the measured state, where closeness is defined in terms of best matching the measurement results. We thus minimize the cost function $\sum_{i=0}^{4^4-1} |p_i - \text{Tr}(\rho_{\text{ph}} \sigma_i)|^2$, subject to $\text{Tr}(\rho_{\text{ph}}) = 1$ and $\rho_{\text{ph}} \geq 0$. We find the optimal $\rho_{\text{ph}}^{\text{opt}}$ using the convex-optimization package *cvxpy* via *cvx-fit* in Qiskit [75]. The fidelity to a target pure state, $|\psi\rangle$, is then computed as

$$F = \langle \psi | \rho_{\text{ph}}^{\text{opt}} | \psi \rangle. \quad (3.13)$$

One can further project ρ_{ph} onto the codespace to obtain a logical state ρ_L using

$$\rho_L = \frac{1}{2} \sum_i \frac{\text{Tr}(\rho_{\text{ph}} \sigma_i^L)}{\text{Tr}(\rho_{\text{ph}} I_L)} \sigma_i^L, \quad \sigma_i^L \in \{I_L, X_L, Y_L, Z_L\} \quad (3.14)$$

where I_L is the projector onto the codespace. Here, we can compute the logical fidelity F_L using Eq. Equation (3.13).

3.4.3 Process tomography in the codespace

A general single-qubit gate can be described [64] by a Pauli transfer matrix (PTM) $\mathcal{R}$ that maps an input state described by $p_i = \langle \sigma_i \rangle$, $\sigma_i \in \{I, X, Y, Z\}$, with $p_0 = 1$, to an output state p' :

$$p'_j = \sum_i \mathcal{R}_{ij} p_i. \quad (3.15)$$

To construct $\mathcal{R}$ in the codespace, we use an overcomplete set of input states,

$$\{|0_L\rangle, |1_L\rangle, |+\rangle, |-\rangle, |+\rangle, |-\rangle\},$$

and their corresponding output states and perform linear inversion. The input and output logical states are characterized using state tomography of the data qubits to find the four-qubit state ρ , which is then projected to the codespace using:

$$\rho_i^L = \frac{\text{Tr}(\rho \sigma_i^L)}{\text{Tr}(\rho I_L)}, \quad \sigma_i^L \in \{I_L, X_L, Y_L, Z_L\}, \quad (3.16)$$

We find that all the measured logical states already satisfy the constraints of a physical density matrix. This is likely to happen as one-qubit states that are not very pure usually lie within the Bloch sphere even within the uncertainty in the measurement. The constructed LPTM, however, might not satisfy the constraints of a physical quantum channel, that is, trace preservation and complete positivity (TPCP). These are better expressed by switching from the PTM representation to the Choi representation. The Choi state $\rho^{\mathcal{R}}$ can be computed as

$$\rho^{\mathcal{R}} = \frac{1}{4} \sum_{i,j} \mathcal{R}_{ij} \sigma_j^T \otimes \sigma_i, \quad (3.17)$$

where the first tensor-product factor corresponds to an auxiliary subsystem. The TPCP constraints are $\text{Tr}(\rho_{\text{ph}}^{\mathcal{R}}) = 1$, $\rho_{\text{ph}}^{\mathcal{R}} \geq 0$ and $\text{Tr}_1(\rho_{\text{ph}}^{\mathcal{R}}) = 1/2$, where Tr_1 is the partial trace over the auxiliary subsystem. In other words, $\rho_{\text{ph}}^{\mathcal{R}}$ is a density matrix satisfying an extra constraint. We then find the optimal $\rho_{\text{ph}}^{\mathcal{R},\text{opt}}$ using the same convex-optimization methods as for state tomography and adding this extra constraint [64, 76]. We compute the corresponding LPTM via

$$(\mathcal{R}_{\text{ph}}^{\text{opt}})_{ij} = \text{Tr}(\rho_{\text{ph}}^{\mathcal{R},\text{opt}} \sigma_j^T \otimes \sigma_i). \quad (3.18)$$

and the average logical gate fidelity using

$$F_L^G = \frac{\text{Tr}(\mathcal{R}_{\text{ideal}}^\dagger \mathcal{R}_{\text{ph}}^{\text{opt}}) + 2}{6}, \quad (3.19)$$

where $\mathcal{R}_{\text{ideal}}$ is the LPTM of the ideal target gate.

3.4.4 Extraction of error-detection rate

The fraction of post-selected data P in the repetitive error detection experiment (Fig. Figure 3.4b) decays exponentially with the number of cycles n . This is consistent with a constant error-detection rate per cycle γ . We extract this rate by fitting the function

$$P(n) = A(1 - \gamma)^n. \quad (3.20)$$

3.5 Supplementary Information

This supplement provides additional information in support of statements and claims made in the previous sections.

3.5.1 Device characteristics

Characteristics of the device and residual ZZ coupling are shown in Tab. 4.1 and Fig. 3.5, respectively.

Qubit	D ₁	D ₂	D ₃	D ₄	A ₁	A ₂	A ₃
Qubit frequency at sweetspot, $\omega_q/2\pi$ (GHz)	6.433	6.253	4.535	4.561	5.770	5.881	5.785
Transmon anharmonicity, $\alpha/2\pi$ (MHz)	-280	—	-320	—	-290	-285	—
Readout frequency, $\omega_r/2\pi$ (GHz)	7.493	7.384	6.913	6.645	7.226	7.058	7.101
Relaxation time, T_1 (τ s)	27	44	32	102	38	58	43
Ramsey dephasing time, T_2^* (τ s)	44	55	51	103	55	60	52
Echo dephasing time, $T_{2,\text{echo}}$ (τ s)	59	70	55	117	69	79	73
Best multiplexed readout fidelity, F_{RO} , (%)	98.6	98.9	96.0	96.5	98.6	94.2	98.9
Single-qubit gate fidelity, F_{SQ} , (%)	99.95	99.86	99.83	99.98	99.95	99.91	99.95

Table 3.2: **Summary of frequency, coherence and readout parameters of the seven transmons.** Coherence times are obtained using standard time-domain measurements [13]. Note that temporal fluctuations of several τ s are typical for these values. The multiplexed readout fidelity, F_{RO} , is the average assignment fidelity [77] extracted from single-shot readout histograms after mitigating residual excitation using initialization by measurement and post-selection [78, 79].

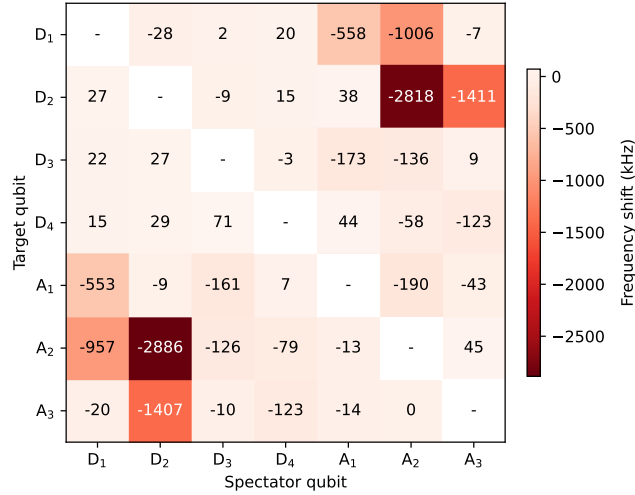


Figure 3.5: **Residual ZZ-coupling matrix.** Measured residual ZZ coupling between all transmon pairs at the bias point (their simultaneous flux sweetspot [74]). Each matrix element denotes the frequency shift that the target qubit experiences due to the spectator qubit being in the excited state, $|1\rangle$. The procedure used for this measurement is similar to the one described in Ref. [80].

3.5.2 Parity-check performance

Parity check assignment fidelity is shown in Fig. 3.6.

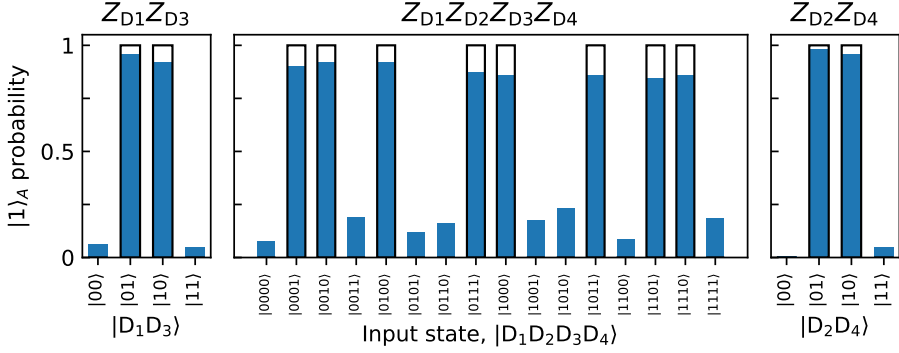


Figure 3.6: **Characterization of the assignment fidelity of Z-type parity checks.** (a) $Z_{D_1}Z_{D_3}$, (b) $^*Z_{D_1}Z_{D_2}Z_{D_3}Z_{D_4}$, and (c) $Z_{D_2}Z_{D_4}$ parity checks implemented using A_1 , A_2 , and A_2 , respectively. Each parity check is benchmarked by preparing the relevant data qubits in a computational state and then measuring the probability of ancilla outcome $m_{A_i} = -1$. Measured (ideal) probabilities are shown as solid blue bars (black wireframe). From the measured probabilities we extract average assignment fidelities 94.2%, 86.1% and 97.2%, respectively. *This parity check implements the $X_{D_1}X_{D_2}X_{D_3}X_{D_4}$ stabilizer measurement with the addition of single-qubit gates on data qubits to perform a change of basis.

3.5.3 Process tomography

Full logical process tomography data is shown in Fig. 3.7.

3.5.4 Logical state stabilization

Logical state stabilization results for states $|0_L\rangle$, $|1_L\rangle$, $|+_L\rangle$ and $|-_L\rangle$ is shown in Fig. 3.8.

Logical error rate

Here, we study the error rate of the logical qubit and compare it to that of a physical qubit. The probability for a logical error on an eigenstate of O_L after n cycles is given by

$$P_{\text{error}}^L = \frac{1 - |\langle O_L \rangle(n)|}{2}. \quad (3.21)$$

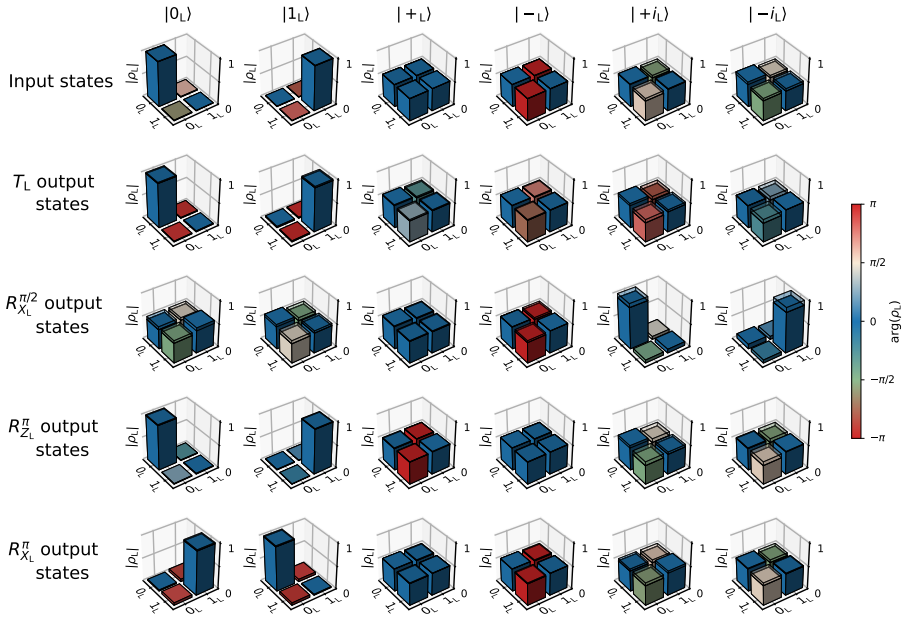


Figure 3.7: **Full set of logical states measured in the logical process tomography procedure.** Measured input and output logical states for each logical gate. Each state is measured using the procedure described in the Methods.

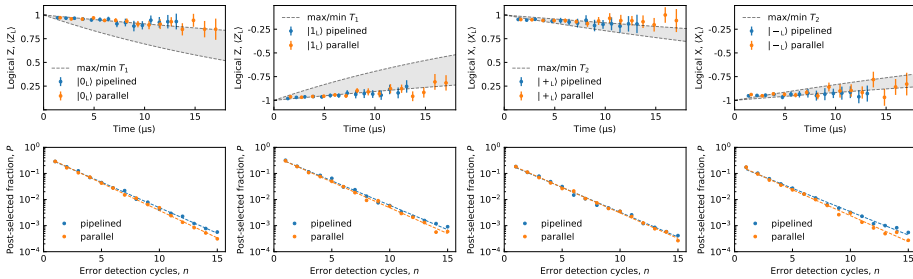


Figure 3.8: **Stabilization of logical cardinal states by repetitive error detection using the pipelined and parallel schemes.** From left to right, the stabilized logical states are $|0_L\rangle$, $|1_L\rangle$, $|+_L\rangle$ and $|-_L\rangle$. For each logical state, the top panel shows the evolution of the relevant logical operator as a function of number of cycles, n , plotted versus wall-clock time. Error bars are estimated based on the statistical uncertainty given by $P(n)$. The shaded area indicates the range of physical qubit T_1 values (a and b) and $T_{2,\text{echo}}$ values (c and d) plotted on the right-axis. Each bottom panel shows the corresponding post-selected fraction of data, $P(n)$.

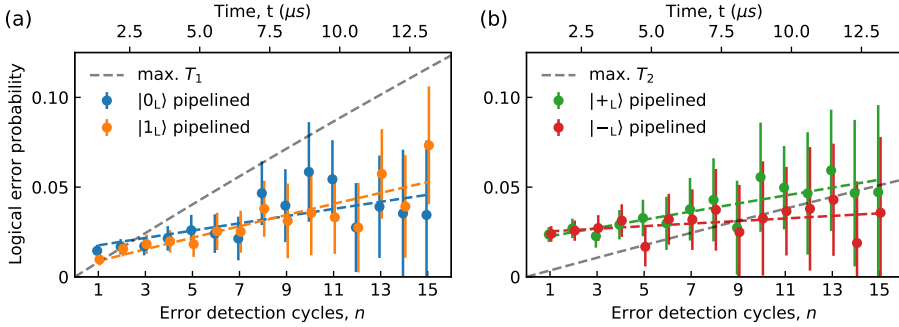


Figure 3.9: Logical error probability versus number of error detection cycles. Logical error probability after n cycles of error detection for states $|0_L\rangle$, $|1_L\rangle$ (a) and $|+_L\rangle$, $|-_L\rangle$ (b) measured using the pipelined scheme. For comparison, the grey dashed curves in (a) and (b) correspond to the physical error probability of the best T_1 and $T_{2,\text{echo}}$ respectively. The logical error rate per round of detection, extracted by fitting the data (colored dashed lines), is 0.43%, 0.67%, 0.49% and 0.15% respectively.

For eigenstates of Z_L we compare the measured error rate to the error experienced due to T_1 on a physical qubit (Fig. 3.9a). In a physical qubit in the excited state, $|1\rangle$, this error is given by

$$P_{\text{error}} = 1 - e^{-t/T_1}. \quad (3.22)$$

For eigenstates of X_L (Fig. 3.9b) we now consider the error experienced due to $T_{2,\text{echo}}$. This error for physical qubits whose state lies on the equator of the Bloch sphere is

$$P_{\text{error}} = \frac{1 - e^{-t/T_{2,\text{echo}}}}{2}. \quad (3.23)$$

We find that the logical error rates for all states, corresponding to the slopes of the colored dashed curves in Fig. 3.9, are lower than the corresponding best physical error rates.

3.5.5 Fault tolerance of logical operations

Fault tolerance of an operation

We begin by elaborating the definition of a fault-tolerant logical operation. We consider a single fault occurring during the circuit implementing the logical operation, where a fault can refer to any single-qubit Pauli error following a single-qubit gate or an idling step, or any two-qubit Pauli error following a two-qubit gate or a measurement error. Furthermore, a single fault can also refer to any single-qubit error on the input state of the logical operation. Thus we consider the performance of the logical operation either when there is a single error at input or a fault in the logical operation. In the context of error detection, the logical operation

such as state initialization or gate execution, is fault-tolerant if any such fault either produces a non-trivial syndrome (in case the circuit involves the measurement of the stabilizers) and is thus post-selected out or leads to an outgoing state that is either the desired logical state or any logical state together with a detectable error. This implies that if the logical operation is followed up by a fictitious and ideal measurement of the stabilizers, the detectable error would lead to a non-trivial syndrome and be post-selected out, ensuring that the outgoing state could only be the desired logical state. For a fault-tolerant logical measurement we require that the logical measurement outcome is correct, i.e. if it is applied to a logical state with single error or a fault happens during the logical measurement, we either post-select or get the correct outcome.

Logical state initialization using stabilizer measurements

We perform fault-tolerant initialization of the logical cardinal states $|0_L\rangle$, $|1_L\rangle$, $|+_L\rangle$ and $|-_L\rangle$. We focus on the initialization of $|0_L\rangle$ and $|1_L\rangle$ and then extend these arguments to $|+_L\rangle$ and $|-_L\rangle$. To prepare $|0_L\rangle$ ($|1_L\rangle$) the data-qubit register is prepared in the state $|0000\rangle$ ($|1010\rangle$) and a single round of stabilizer measurements is performed. Consider any single-qubit error occurring during the initialization of the data qubits. Any such error will be detected by the following stabilizer measurement and post-selected out. Then, consider a fault occurring during the stabilizer measurement, implemented by the circuits shown in Fig. 4. Any single-qubit error on the data qubits following an idling step either leads to non-trivial syndrome (if it occurs before the two-qubit gate) or constitutes a detectable error (if it occurs after the two-qubit gate). A single-qubit error following any of the single-qubit gates will similarly either produce non-trivial syndrome or lead to an error that is either detectable or an element in $\mathcal{S}$. A measurement error on ancilla qubits A_1 or A_3 will always produce a non-trivial syndrome (since the input state is an eigenstate of the measured Z type stabilizers) and will thus be post-selected out. However, a measurement error on A_2 can still lead to trivial syndrome (as the input state is not an eigenstate of the X -type stabilizer). This will result in the preparation of the desired logical state together with a detectable phase-flip error on any of the qubits. Any two-qubit Pauli error after each CZ gate involved in the measurement of the Z -type stabilizers will lead to the preparation of the desired logical state together with an error that is either in $\mathcal{S}$ (for example in the case when a bit-flip error occurs on the ancilla qubit and phase-flip error on the data-qubit following the first CZ gate of the circuits) or one that is detectable. The same statement holds for any two-qubit error after each of the CZ gates involved in the X type stabilizer check. Here the order of the gates is crucial to ensure that any two-qubit error after the second CZ gate of the circuit is detectable [65]. The fault-tolerant preparation of $|+_L\rangle$ ($|-_L\rangle$) involves initializing the data-qubit register in the state $|++++\rangle$ ($|++--\rangle$) instead. The arguments for the fault-tolerance of this operation follow closely the ones presented for $|0_L\rangle$ ($|1_L\rangle$) with the only difference being that a single measurement error on ancilla qubits A_1 or A_3 can now lead to a trivial outcome and an outgoing state that involves a detectable (bit-flip)

error, while a measurement error on A_2 will instead always lead to a non-trivial syndrome. When initializing arbitrary logical states by preparing the data qubit register in the state given by Eq. 6, the procedure is fault-tolerant only when preparing $|0_L\rangle$ (corresponding to $\theta = 0$ and $\phi = 0$) or $|1_L\rangle$ (corresponding to $\theta = \pi$ and $\phi = 0$), which are discussed above. When preparing $|+_L\rangle$ (corresponding to $\theta = \pi/2$ and $\phi = 0$) or $|-_L\rangle$ (corresponding to $\theta = \pi/2$ and $\phi = \pi$), the input states are $|+0+0\rangle$ and $|+0-0\rangle$ respectively. In these cases a single fault (for example a phase-flip error on qubit D_3 on the input state) is not detectable by the stabilizer measurement and instead leads to the initialization of (the opposite states) $|-_L\rangle$ and $|+_L\rangle$, respectively. The preparation of any other state on the equator of the Bloch sphere is not fault-tolerant either, following the same reasoning: an under- or overrotation of ϕ will directly translate to an error at the logical level.

Fault-tolerant logical measurements

We now consider the fault-tolerance of the logical measurement, which is performed following the procedure described in the Results (see Fig. 2). The only fault to consider in these circuits is a measurement error on one of the data qubits. When measuring X_L or Z_L , any such error will result in a non-trivial syndrome (given the assumption that the input state is in the codespace) and the logical measurement outcome is post-selected out. When the fault is instead a single-qubit error on the input state, bringing this state outside of the codespace, the fault-free logical measurement will either detect this error or this error will not have an affect on the logical measurement outcome. For the non-fault-tolerant measurement of Y_L only the value for the $Z_{D2}Z_{D4}$ stabilizer can be computed and used to detect errors on D_2 and D_4 . Thus a single fault (for example a measurement error on either D_1 or D_3) can lead to an incorrect logical measurement outcome, making this operation non-fault-tolerant.

Transversal logical gates and non-fault-tolerant gate injection

The logical gates $R_{X_L}^\pi$ and $R_{Z_L}^\pi$ (shown in Fig. 3d) are clearly fault-tolerant as any single-qubit error following any of the single-qubit gates involved in the circuits is detectable. At the same time the transversal execution of these gates ensures that no single qubit error on the input state can spread to two or more qubits, ensuring that any such fault is detectable. These fault-tolerant properties do not hold when we consider the T_L and $R_{X_L}^{\pi/2}$ logical gates implemented by the gate-by-measurement circuits shown in Fig. 3b (and Fig. 3d). For example a bit-flip error on A_2 following the first single-qubit gate of the circuit will result in a logical error. More generally, any under- or over-rotation in the rotation angle θ used in preparing the ancilla qubit in $|A_\theta\rangle$ translates to a different rotation at the logical level than desired.

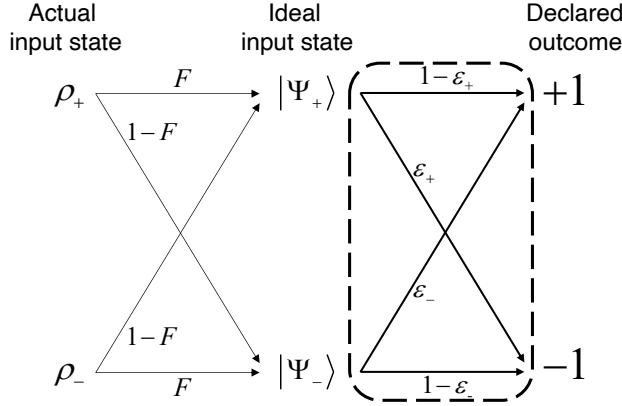


Figure 3.10: **Probability flow diagram for physical qubit readout.** Please see text for the definition of all variables shown. The characterization of physical qubit readout robust to initialization errors determines the probabilities within the dashed box.

3.5.6 Quantifying the logical assignment fidelity

We start this section reviewing how the readout fidelity F^R of a physical qubit is standardly quantified from the contrast of a Rabi oscillation when the input states $\rho_{\pm}$ closest to the eigenstates $|\Psi_{\pm}\rangle$ of the measured observable O have limited fidelity F (assumed equal for both). Evidently, we want F^R to quantify the performance of readout only, independent of errors in the input state. To this end, consider the probability flow diagram of Fig. 3.10. We define F^R as the average probability of proper assignment for perfect input states $|\Psi_{\pm}\rangle$. Therefore, $F^R = 1 - (\epsilon_+ + \epsilon_-)/2$, where $\epsilon_{\pm}$ is the probability of wrongly assigning outcome ∓ 1 for input state $|\Psi_{\pm}\rangle$. The positive and negative extremes of the Rabi oscillation are

$$\langle O \rangle_{\max} = F\bar{\epsilon}_+ + \bar{F}\epsilon_- - F\epsilon_+ - \bar{F}\bar{\epsilon}_-. \quad (3.24)$$

$$\langle O \rangle_{\min} = -F\bar{\epsilon}_- - \bar{F}\epsilon_+ + F\epsilon_- + \bar{F}\bar{\epsilon}_+, \quad (3.25)$$

where $\bar{F} = 1 - F$ and $\bar{\epsilon}_{\pm} = 1 - \epsilon_{\pm}$. Combining these expressions and simplifying terms, it follows that the contrast of the Rabi oscillation (defined as half the peak-to-peak range), is

$$\frac{\langle O \rangle_{\max} - \langle O \rangle_{\min}}{2} = (2F - 1)(2F^R - 1). \quad (3.26)$$

Turning over to the logical qubit, we similarly define the logical readout fidelity F_L^R as the average probability of proper assignment for perfectly prepared logical states $|\Psi_{L\pm}\rangle$ i.e., the logical states that are eigenstates of the measured observable O_L with eigenvalue ± 1 . To this end, it is tempting to apply the above equation to the oscillations in Fig. 2, simply substituting $F \rightarrow F_L$ and $F^R \rightarrow F_L^R$. However, this approach is not accurate. This is because the

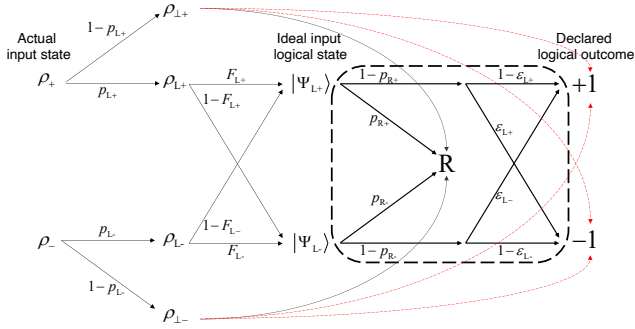


Figure 3.11: **Probability flow diagram for logical readout.** Please see text for the definition of all variables shown. The characterization of logical qubit readout robust to initialization errors determines the probabilities within the dashed box. The probability for states outside the codespace to not be rejected (red dashed curves) is the primary reason why using Eq. 3.26 with the substitutions $F \rightarrow F_L$ and $F^R \rightarrow F_L^R$ does not yield an accurate estimate of F_L^R .

probability flow diagram for the logical qubit, shown in Fig. 3.11, is more complex. The quantities we seek to determine are those inside the dashed box, describing logical readout on perfect input logical states: $p_{R\pm}$ is the probability that the experimental logical measurement on $|\Psi_{L\pm}\rangle$ is rejected (R), which occurs whenever the data-qubit outcome string produces a value of -1 on at least one of the stabilizers computable from the string; $\epsilon_{L\pm}$ is the probability of wrongly assigning logical outcome ∓ 1 for $|\Psi_{L\pm}\rangle$, conditioned on no rejection. Using these definitions, $F_L^R = 1 - (\epsilon_{L+} + \epsilon_{L-})/2$. Outside the dashed box, $\rho_{\pm}$ is the experimental input state closest to $|\Psi_{L\pm}\rangle$. This imperfect input state has probability $p_{L\pm}$ of being in the codespace and its projection onto the codespace, $\rho_{L\pm}$, has fidelity $F_{L\pm}$ to $|\Psi_{L\pm}\rangle$. Finally, $\rho_{\perp\pm}$ is the projection of $\rho_{\pm}$ outside the codespace.

We now discuss the more accurate method used to quantify F_L^R that does not rely on Eq. 3.26. We first consider the transformation of $|\Psi_{L\pm}\rangle$ by the pre-rotations that we perform when measuring in each cardinal logical basis, $O_L \in \{Z_L, X_L, Y_L\}$. For Z_L there are no measurement pre-rotations, so the states are

$$|0_L\rangle = \frac{1}{\sqrt{2}} (|0000\rangle + |1111\rangle), \quad (3.27)$$

$$|1_L\rangle = \frac{1}{\sqrt{2}} (|0101\rangle + |1010\rangle). \quad (3.28)$$

For X_L , the transformed states are

$$R_{1y}^{-\pi/2} R_{2y}^{-\pi/2} R_{3y}^{-\pi/2} R_{4y}^{-\pi/2} |+_L\rangle = \frac{1}{2} (|0000\rangle + |0101\rangle + |1010\rangle + |1111\rangle), \quad (3.29)$$

$$R_{1y}^{-\pi/2} R_{2y}^{-\pi/2} R_{3y}^{-\pi/2} R_{4y}^{-\pi/2} |-_L\rangle = \frac{1}{2} (|0011\rangle + |0110\rangle + |1001\rangle + |1100\rangle). \quad (3.30)$$

Finally, for Y_L , these are

$$R_{1x}^{\pi/2} R_{3y}^{-\pi/2} |+_L\rangle = \frac{1}{2} (|0000\rangle - i|0111\rangle + i|1010\rangle + |1101\rangle), \quad (3.31)$$

$$R_{1x}^{\pi/2} R_{3y}^{-\pi/2} |-_L\rangle = \frac{1}{2} (-|0010\rangle - i|0101\rangle - i|1000\rangle + |1111\rangle). \quad (3.32)$$

The above expressions make clear, for each logical cardinal basis, which data-qubit outcome strings are rejected and which ones are accepted with declared logical outcome $+1$ or -1 . For completeness, all cases are detailed in Table 3.3.

The key experimental input needed to proceed is the data-qubit assignment probability matrix A , shown in Fig. 3.12. Each element of this 16×16 matrix gives the experimental probability of measuring a string of data outcomes $(m_{D1}, m_{D2}, m_{D3}, m_{D4})$, $m_{Di} \in \{-1, 1\}$ (varying across rows) when performing simultaneous readout of the data qubits having prepared them in physical computational state $|n_{D1} n_{D2} n_{D3} n_{D4}\rangle$, $n_{Di} \in \{0, 1\}$ (varying across columns). These computational states are prepared by applying the needed parallel combination of R_{ix}^π pulses on data qubits starting with all qubits (including ancillas) initialized in $|0\rangle$. With A in hand, it is straightforward to compute the probabilities for all strings of data-qubit outcomes for each choice of O_L and $|\Psi_{L\pm}\rangle$. This is given by $A\vec{p}$, where $\vec{p}$ is vector (size 16) whose elements are the probabilities (in the physical data-qubit computational basis) of the corresponding state in Eqs. 3.27-3.32. For example, $\vec{p} = (1/2, 0, \dots, 0, 1/2)^T$ for Z_L and $|0_L\rangle$, and $\vec{p} = (1/4, 0, 0, 0, 0, 1/4, 0, 0, 0, 0, 1/4, 0, 0, 0, 0, 1/4)^T$ for X_L and $|+_L\rangle$. From $A\vec{p}$ and the rejection and logical assignment rules in Table 3.3, it is straightforward to compute all the probabilities within the dashed box of Fig. 3.11. The final results are presented in Table 3.4. The key assumption behind this analysis is that errors induced by single-qubit gates (both during preparation of the physical data-qubit computational states needed for determination of A and the measurement pre-rotations when performing logical measurement in X_L and Y_L) are small compared to the errors induced by data-qubit readout. This assumption is safe given the performance metrics summarized in Table 4.1.

3.5.7 Numerical analysis

Leakage in experiment

We observe a clear signature of leakage accumulation with the increasing number of error-detection cycles in the single-shot readout histograms obtained at the end of each experiment. In Fig. 3.13 we show examples of this accumulation for D_2 , D_3 and A_2 at cycles $n = 1$, $n = 8$ and $n = 15$. For dispersive readout, a transmon in state $|2\rangle$ induces a different frequency shift in the readout resonator compared to state $|0\rangle$ or $|1\rangle$. The increased number of data points at $n = 8$ and $n = 15$ shown in Fig. 3.13, following a Gaussian distribution with a mean and standard deviation different from those observed at $n = 1$ is thus a clear manifestation of leakage to the higher-excited states (mostly to $|2\rangle$). We believe that the dominant source of leakage in our processor are the CZ gates [9, 11]. However, the leakage rate L_1 for each gate

Data-qubit outcome ($m_{D1}, m_{D2}, m_{D3}, m_{D4}$)	Logical assignment			Data-qubit outcome ($m_{D1}, m_{D2}, m_{D3}, m_{D4}$)	Logical assignment		
	Z_L	X_L	Y_L		Z_L	X_L	Y_L
(+1, +1, +1, +1)	+1	+1	+1	(-1, +1, +1, +1)	$R(\mathcal{S}_1)$	$R(\mathcal{S}_2)$	-1
(+1, +1, +1, -1)	$R(\mathcal{S}_3)$	$R(\mathcal{S}_2)$	$R(\mathcal{S}_3)$	(-1, +1, +1, -1)	$R(\mathcal{S}_1, \mathcal{S}_3)$	-1	$R(\mathcal{S}_3)$
(+1, +1, -1, +1)	$R(\mathcal{S}_1)$	$R(\mathcal{S}_2)$	-1	(-1, +1, -1, +1)	-1	+1	+1
(+1, +1, -1, -1)	$R(\mathcal{S}_1, \mathcal{S}_3)$	-1	$R(\mathcal{S}_3)$	(-1, +1, -1, -1)	$R(\mathcal{S}_3)$	$R(\mathcal{S}_2)$	$R(\mathcal{S}_3)$
(+1, -1, +1, +1)	$R(\mathcal{S}_3)$	$R(\mathcal{S}_2)$	$R(\mathcal{S}_3)$	(-1, -1, +1, +1)	$R(\mathcal{S}_1, \mathcal{S}_3)$	-1	$R(\mathcal{S}_3)$
(+1, -1, +1, -1)	-1	+1	-1	(-1, -1, +1, -1)	$R(\mathcal{S}_1)$	$R(\mathcal{S}_2)$	+1
(+1, -1, -1, +1)	$R(\mathcal{S}_1, \mathcal{S}_3)$	-1	$R(\mathcal{S}_3)$	(-1, -1, -1, +1)	$R(\mathcal{S}_3)$	$R(\mathcal{S}_2)$	$R(\mathcal{S}_3)$
(+1, -1, -1, -1)	$R(\mathcal{S}_1)$	$R(\mathcal{S}_2)$	+1	(-1, -1, -1, -1)	+1	+1	-1

Table 3.3: **Logical assignments from data-qubit measurement outcomes.** When measuring in the specified logical cardinal basis (as shown in Fig. 2a), the final string of data-qubit outcomes is rejected (R) whenever at least one of the computable stabilizers has value -1 (indicated within parentheses). The computable stabilizers are $\mathcal{S}_1$ and $\mathcal{S}_3$ when measuring Z_L , $\mathcal{S}_2$ when measuring X_L , and $\mathcal{S}_3$ when measuring Y_L . When the data-qubit outcome string is accepted, the logical assignment (+1 or -1) is given by the appropriate product of data-qubit outcomes: $m_{D1} \times m_{D2}$ for Z_L , $m_{D1} \times m_{D3}$ for X_L , and $m_{D1} \times m_{D2} \times m_{D3}$ for Y_L .

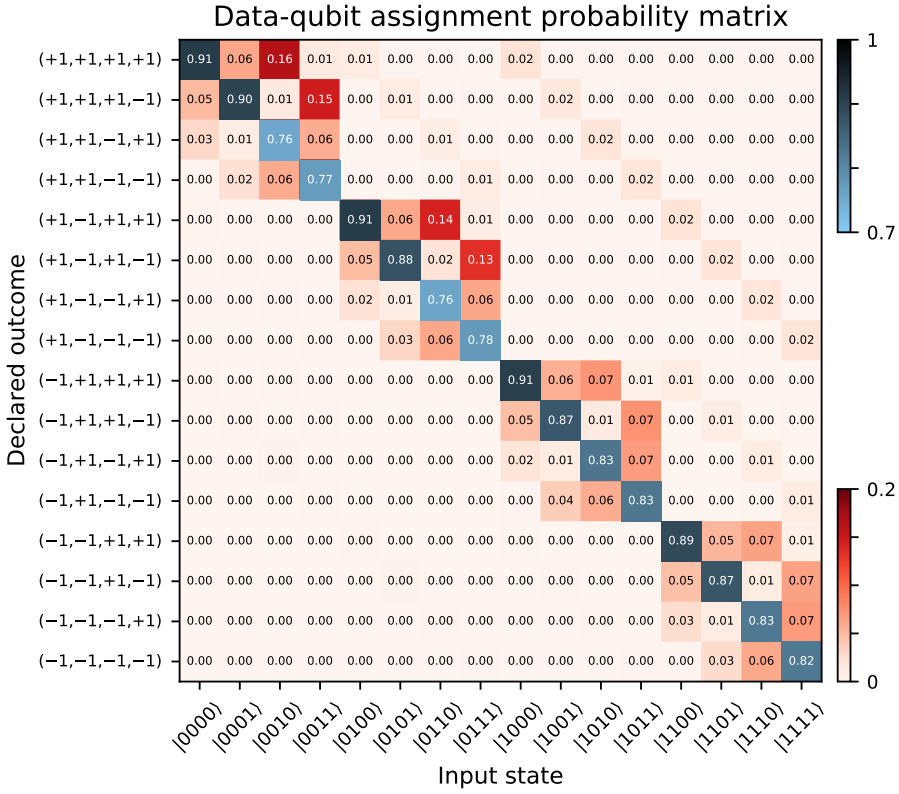


Figure 3.12: **Experimental data-qubit assignment probability matrix.** Each element of A gives the experimental probability of measuring outcome string $(m_{D1}, m_{D2}, m_{D3}, m_{D4})$ (varying across rows) when performing simultaneous measurement of the data qubits prepared in $|n_{D1}n_{D2}n_{D3}n_{D4}\rangle$, $n_{Di} \in \{0, 1\}$ (varying across columns).

has not been experimentally characterized, e.g., by performing leakage-modified randomized benchmarking experiments [81, 82]. This is because our CZ tune-up procedure is performed in a parity-check block unit. This maximizes the performance of the parity-check but makes the gate unfit for randomized benchmarking protocols. We can estimate the population $p^{\mathcal{L}}(n)$ in the leakage subspace $\mathcal{L}$ at cycle n from the single-shot readout histograms. We perform a fit of a triple Gaussian model to the histograms from which we extract the voltage that allows for the best discrimination of $|2\rangle$ from $|1\rangle$ and $|0\rangle$. The leaked population $p^{\mathcal{L}}(n)$ is then given by the fraction of shots declared as $|2\rangle$ over the total number of shots. Assuming that leakage is only induced by the CZ gates (on the transmon being fluxed to perform the gate) and that each CZ gate has the same leakage rate L_1 , we can use the Markovian model presented in Ref. [70] to estimate the L_1 value leading to the observed population $p^{\mathcal{L}}(n)$. This analysis

Logical measurement basis O_L	Z_L	X_L	Y_L
$ \Psi_{L+}\rangle$	$ 0_L\rangle$	$ +_L\rangle$	$ +i_L\rangle$
$ \Psi_{L-}\rangle$	$ 1_L\rangle$	$ -_L\rangle$	$ -i_L\rangle$
P_{R+}	0.184	0.164	0.078
P_{R-}	0.170	0.118	0.073
ϵ_{L+}	0.003	0.016	0.089
ϵ_{L-}	0.002	0.010	0.083
F_L^R	0.998	0.987	0.914

Table 3.4: **Quantified performance of logical measurement.** Final results of the analysis performed to quantify logical measurement in the logical cardinal bases without corruption from initialization errors. See Fig. 3.11 for reference. The extracted logical readout fidelities are those quoted in the main text.

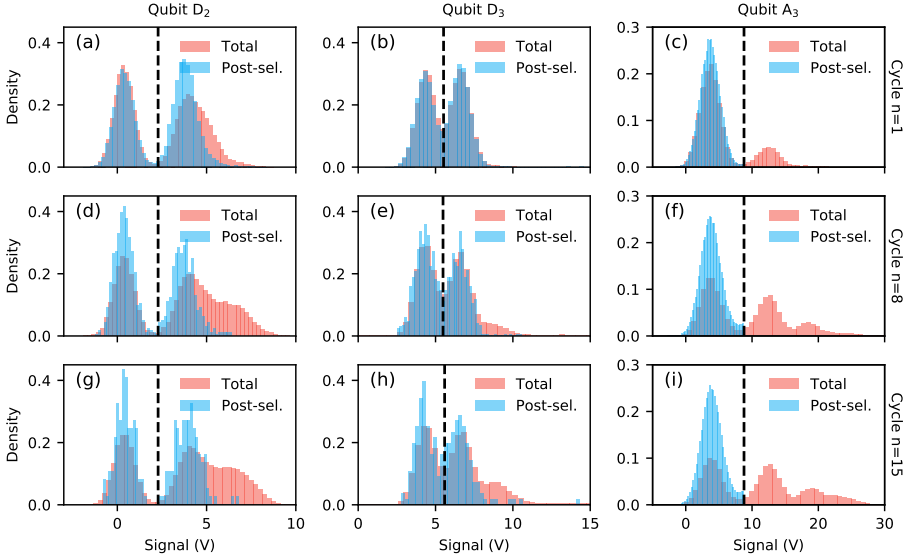


Figure 3.13: **Signature of transmon leakage in experimental data.** Single-shot readout histograms obtained at cycle n over all shots (red) and the post-selected shots based on detecting no error in any cycles up to n (blue) for D_2 (left), D_3 (middle) and A_2 (right) and at cycle $n = 1$ (top row), $n = 8$ (middle row) and $n = 15$ (bottom row). The dashed black lines indicate the thresholds used to discriminate $|0\rangle$ from $|1\rangle$.

gives a L_1 estimate in the approximate range $1 - 4\%$ for most transmons. However, we do not consider these estimates to be accurate due to the low fidelity with which $|2\rangle$ can

be distinguished from $|1\rangle$ and instead treat L_1 as a free parameter in our simulations (see below).

The histograms of the post-selected shots in Fig. 3.13 demonstrate that post-selection rejects runs where leakage on those transmons occurred. Thus, while leakage may considerably impact the error-detection rate in the experiment [70], we do not expect it to significantly affect the fidelity of the logical initialization, and gates.

Density-matrix simulations

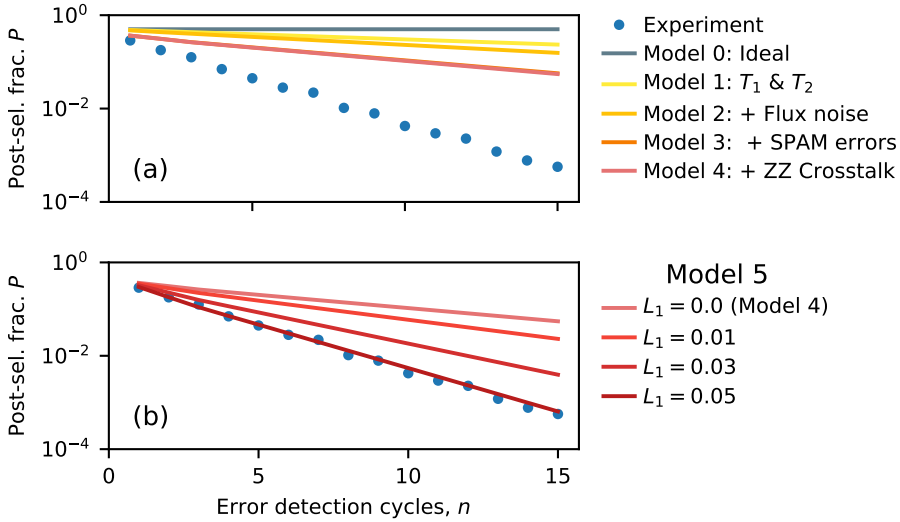


Figure 3.14: **Simulation of error-detection rate.** Post-selected fraction P as a function of the number n of error-detection cycles for $|0_L\rangle$. The experimental P (blue dots) is compared to numerical simulation under various models (solid curves). (a) Simulated P obtained by incremental addition of error sources starting from the no-error (Model 0, gray); qubit relaxation and dephasing (Model 1, yellow); extra dephasing due to flux noise away from the sweetspot (Model 2, amber); state preparation and measurement errors (Model 3, orange); and crosstalk due to residual ZZ interactions (Model 4, red). (b) Simulated P for Model 5 adding CZ gate leakage with 4 different values of L_1 , the leakage per CZ gate, assumed equal for all CZ gates.

We perform numerical density-matrix simulations using the *quantumsim* package [83] to study the impact of the expected error sources on the performance of the code. We focus on repetitive error detection using the pipelined scheme and with the logical qubit initialized in $|0_L\rangle$. In Fig. 3.14a, we show the post-selected fraction $P(n)$ as a function of the number n of error-detection cycles for a series of models. Model 0 is a no-error model, which we take

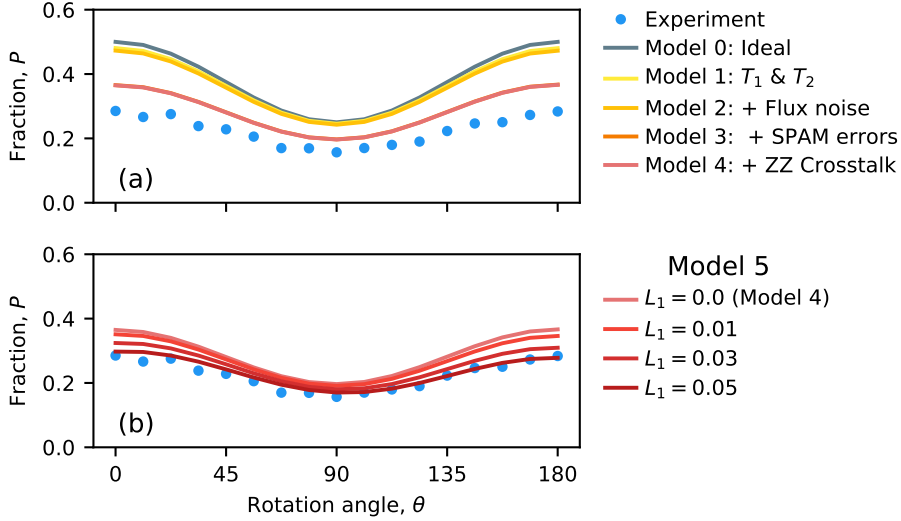


Figure 3.15: **Simulated post-selected fraction.** Post-selected fraction P of Fig. 2f for the Z_L measurement with the same error models used in Fig. 3.14.

as the starting point of the comparison. Model 1 adds amplitude and phase damping experienced by the transmon. Model 2 adds the increased dephasing away from the sweetspot arising from flux noise. Model 3 adds residual qubit excitation and readout (SPAM) errors. Finally, Model 4 adds crosstalk due to the residual ZZ coupling during the coherent operations of the stabilizer measurement circuits. The details of each model and their input parameters drawn from experiment are detailed below. We find that the dominant contributors to the error-detection rate are SPAM errors and decoherence. However, we also observe that the noise sources included through Model 4 clearly fail to quantitatively capture the decay of the post-selected fraction observed in experiment.

We believe that an important factor behind the observed discrepancy is the presence of leakage, as suggested by the single-shot readout histograms in Fig. 3.13. We consider the leakage per CZ gate L_1 as a free parameter and assume the same value for all CZ gates. We simulate the post-selected fraction for a range of L_1 values, shown in Fig. 3.14b. We observe that $L_1 \approx 5\%$ produces a good match with experiment, suggesting that leakage significantly impacts the error-detection rate observed. We perform a similar analysis now considering the logical measurement of Z_L experiment depicted in Fig. 2f which also finds similar agreement with experimental data (Fig. 3.15). This value of L_1 is significantly higher than achieved in Ref. [11], which used the same device. However, note that in this earlier experiment CZ gates were characterized while keeping all other qubits in $|0\rangle$. Spectator transmons with residual ZZ coupling to either of the transmons involved in a CZ gate can increase L_1 when not

in $|0\rangle$ (which is certainly the case in the present experiment). Note that leakage may also be further induced by the measurement [18], an effect that we do not consider in our simulation. However, the assumption that all CZ gates have the same L_1 , the approximations used in our models, and other error sources that we have not considered here may lead to an overestimation of the true L_1 .

Leakage is an important error source to consider in quantum error correction experiments of larger distance codes, requiring either post-selection based on detection [70] or the use of leakage reduction units [72]. We leave the detailed investigation of the exact leakage rates in our experiment and the mechanisms leading to them to future work.

Error models

Lastly, we detail the error models used in the numerical simulations in Fig. 3.14.

- **Model 1:** We take into account transmons decoherence by including an amplitude-damping channel parameterized by the measured relaxation time T_1 and a phase-damping channel parameterized by the pure-dephasing time at the sweetspot

$$\frac{1}{T_\phi^{\max}} = \frac{1}{T_{2,\text{echo}}} - \frac{1}{2T_1},$$

where $T_{2,\text{echo}}$ is the measured echo dephasing time (see Table 4.1). The qutrit Kraus operators defining these channels are detailed in Ref. [70] and we similarly introduce these channels during idling periods and symmetrically around each single-qubit or two-qubit gate (each period lasting half the duration of the gate).

- **Model 2:** We consider the pure-dephasing rate $1/T_\phi = 2\pi\sqrt{\ln 2}AD_\phi + 1/T_\phi^{\max}$ away from the sweetspot due to the fast-frequency components of the $1/f$ flux noise, where D_ϕ is the flux sensitivity at a given qubit frequency and A is the scaling parameter for the flux-noise spectral density. We use a $\sqrt{A} \approx 3 \mu\Phi_0$, the average of the extracted $\sqrt{A}$ values for D_3 , A_1 and A_2 obtained by fitting the measured decrease of $T_{2,\text{echo}}$ as a function of the applied flux bias, following the model described above. This allows us to estimate the dephasing time at the CZ interaction and parking frequencies, which then parameterize the applied amplitude-phase damping channel inserted during those operations [70]. We neglect the slow-frequency components of the flux noise due to the use of sudden Net Zero pulses, which echo out this noise to first order [9, 11].
- **Model 3:** We further include state-preparation and measurement errors. We consider residual qubit excitations, where instead of the transmon being initialized in $|0\rangle$ at the start of the experiment, it is instead excited to $|1\rangle$ with probability p_e . We extract p_e for each transmon from a double-Gaussian fit to the histogram of the single-shot readout voltages with the transmon nominally initialized in $|0\rangle$ [79]. We model measurement errors via the POVM operators $M_i = \sum_{j=0}^2 \sqrt{P(i|j)} |j\rangle\langle j|$ for $i \in 0, 1, 2$ being

the measurement outcome, while $P(i|j)$ is the probability of measuring the qubit in state $|i\rangle$ when having prepared state $|j\rangle$. We extract the probability $P(Q = |i\rangle) = \text{Tr}(M_i^\dagger M_i \rho)$ of measuring qubit Q in state $|i\rangle$ from simulation, where ρ is the density matrix, while application of the POVM transforms $\rho \rightarrow M_i \rho M_i^\dagger / P(Q = |i\rangle)$. In our simulations we condition on the detection of no error and thus we calculate $P(Q = |0\rangle)$ and then apply M_0 to the state ρ . We obtain $P(0|j)$ for $j \in 0, 1$ from the experimental assignment fidelity matrix [41] (where a heralded initialization protocol was used to prepare the qubits in $|0\rangle$ [78]) and we assume $P(0|2) = 0$, consistent with the observed histograms in Fig. 3.13. At the end of each experiment with n error-detection cycles we calculate the probability P_n^f of obtaining trivial syndromes from the final measurements of the data qubits (see Results). From this and from the probability $P_n(A_i = |0\rangle)$ of measuring ancilla A_i in $|0\rangle$ at cycle n , we calculate the post-selected fraction of experiments defined as $P(n) = P_n^f \prod_n \prod_{i=1}^3 P_n(A_i = |0\rangle)$.

- **Model 4:** We consider the crosstalk due to residual ZZ interactions between transmons. The CZ gates involved in a parity check are jointly calibrated to minimize phase errors for the whole check as one block (see Fig. 3.6). Instead of modeling this crosstalk as an always-on interaction and taking into account the details of the check calibration, we instead capture the net effect of this noise by including it as single-qubit and two-qubit phase errors in each CZ gate. This assumes that the crosstalk only occurs between transmons that are directly coupled, which the measured frequency shifts observed in Fig. 3.5 validate. We characterize the phases picked up during the CZ gates using $k \times 2^{k-1}$ Ramsey experiments for a check involving a total of k transmons (including the ancilla). In each experiment, we perform a Ramsey experiment on one transmon labelled Q_k . Q_k is initialized in a maximal superposition using a $R_x^{-\pi/2}$ pulse, while the remaining $k - 1$ transmons are prepared in each of the 2^{k-1} computational states $|I\rangle$. Following this initialization, the parity check is performed, followed by a rotation of $R_\phi^{-\pi/2}$ (while the other transmons are rotated back to $|0\rangle$) and by a measurement of Q_k . By varying the axis of rotation ϕ , we extract the phase $\phi_{\text{Ram}}^k(I)$ picked up by Q_k with the remaining transmons in state $|I\rangle$. We perform this procedure for each of the k transmons of the check, resulting in a total of $k \times 2^{k-1}$ measured phases, which are arranged in a column vector $\vec{\phi}^{\text{Ram}}$. We parameterize each CZ gate used in the parity check by a matrix $\text{diag}(1, e^{i\phi_{01}}, e^{i\phi_{10}}, e^{i\phi_{11}})$. The column vector $\vec{\phi}^{\text{CZ}}$ then contains all of the phases parameterizing each of the $k - 1$ CZ gates involved in the parity checks, with $k = 3$ for the $Z_{D1}Z_{D3}$ and $Z_{D2}Z_{D4}$ checks and $k = 5$ for the $Z_{D1}Z_{D2}Z_{D3}Z_{D4}$ check. We can express each of the measured phases in the Ramsey experiment as a linear combination of the acquired phases as a result of the CZ interactions between transmons, i.e., $\vec{\phi}^{\text{Ram}} = A \vec{\phi}^{\text{CZ}}$, where the matrix A encodes

the linear dependence. Given the measured $\bar{\phi}^{\text{Ram}}$ we perform an optimization to find the closest $\bar{\phi}^{\text{CZ}}$ as given by

$$\begin{aligned} \min_{\bar{\phi}^{\text{CZ}}} \quad & \sum_i \left(\sum_j A_{ij} \bar{\phi}_j^{\text{CZ}} - \bar{\phi}_i^{\text{Ram}} \right)^2, \\ \text{subject to} \quad & 0 \leq \bar{\phi}_j^{\text{CZ}} < 2\pi. \end{aligned}$$

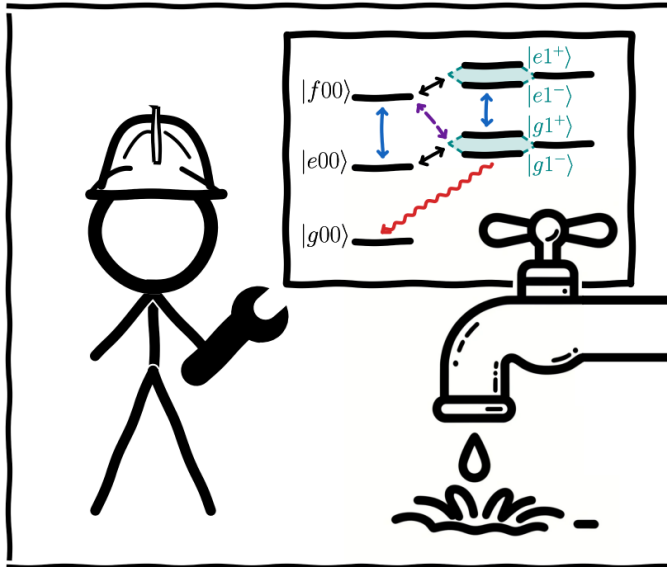
The optimal $\bar{\phi}^{\text{CZ}}$ then captures the net effect of the ZZ crosstalk during the parity checks, which we include in the simulation. We do not model phase errors accrued during the ancilla readout, since in our simulation we condition on each ancilla being measured in $|0\rangle$.

- **Model 5:** We model leakage due to CZ gates following the model and numerical implementation presented in Ref. [70]. Here, we do not consider the phases picked up when non-leaked transmons interact with leaked ones (the leakage-conditional phases [70]) and we set them to their ideal values. We also neglect higher-order leakage effects, such as excitation to higher-excited states or leakage mobility. Thus, we only consider the exchange of population between $|11\rangle$ and $|02\rangle$ given by $4L_1$, except for the CZ between A_1 and D_3 , where the population is instead exchanged with $|20\rangle$ as we use the $|11\rangle$ - $|20\rangle$ avoided crossing for this gate [11].

There remain several relevant error sources beyond those included in our numerical simulation. For example, we do not include dephasing of data or other ancilla qubits induced by ancilla measurement, which we expect to be a relevant error source for comparing the performance of the pipelined and parallel schemes. Also, we only consider the net effect of crosstalk due to residual ZZ interactions during coherent operations of the parity-check circuits, which we include via errors in the single-qubit and two-qubit phases in the CZ gates. Thus, we do not capture the crosstalk present whenever an ancilla is projected to state $|1\rangle$ by the readout but declared to be in $|0\rangle$ instead. Furthermore, as ZZ crosstalk does not commute with the amplitude damping included during the execution of the circuit, we are not capturing the increased phase error rate that this leads to.

ALL-MICROWAVE LEAKAGE REDUCTION UNITS FOR QUANTUM ERROR CORRECTION WITH SUPERCONDUCTING QUBITS

J. F. Marques, H. Ali, B. M. Varbanov, M. Finkel, H. M. Veen, S. L. van der Meer, S. Valles-Sanclemente, N. Muthusubramanian, M. Beekman, N. Haider, B. M. Terhal, L. DiCarlo



Parts of this chapter have been published in *PRL* **130**, 250602 (2023).

4.1 Introduction

Superconducting qubits, such as the transmon [3], are many-level systems in which a qubit is represented by the two lowest-energy states $|g\rangle$ and $|e\rangle$. However, leakage to non-computational states is a risk for all quantum operations, including single-qubit gates [23], two-qubit gates [7, 11, 16] and measurement [18, 19]. While the typical probability of leakage per operation may pale in comparison to conventional qubit errors induced by control errors and decoherence [11, 46], unmitigated leakage can build up with increasing circuit depth. A prominent example is multi-round quantum error correction (QEC) with stabilizer codes such as the surface code [62]. In the absence of leakage, such codes successfully discretize all qubit errors into Pauli errors through the measurement of stabilizer operators [49, 84], and these Pauli errors can be detected and corrected (or kept track of) using a decoder. However, leakage errors fall outside the qubit subspace and are not immediately correctable [85–87]. The signature of leakage on the stabilizer syndrome is often not straightforward, hampering the ability to detect and correct it [58, 70]. Additionally, the build-up of leakage over QEC rounds accelerates the destruction of the logical information [46, 71]. Therefore, despite having low probability per operation, methods to reduce leakage must be employed when performing experimental QEC with multi-level systems.

Physical implementations of QEC codes [38, 39, 88–91] use qubits for two distinct functions: Data qubits store the logical information and, together, comprise the encoded logical qubits. Ancilla qubits perform indirect measurement of the stabilizer operators. Handling leakage in ancilla qubits is relatively straightforward as they are measured in every QEC cycle. This allows for the use of reset protocols [71, 92] without the loss of logical information. Leakage events can also be directly detected using three- or higher-level readout [38] and reset using feedback [57, 78]. In contrast, handling data-qubit leakage requires a subtle approach as it cannot be reset nor directly measured without loss of information or added circuit complexity [93–95]. A promising solution is to interleave QEC cycles with operations that induce seepage without disturbing the qubit subspace, known as leakage reduction units (LRUs) [72, 85, 86, 93, 94, 96–99]. An ideal LRU returns leakage back to the qubit subspace, converting it into Pauli errors which can be detected and corrected, while leaving qubit states undisturbed. By converting leakage into conventional errors, LRUs enable a moderately high physical noise threshold, below which the logical error rate decreases exponentially with the code distance [86, 94]. A more powerful operation called 'heralded leakage reduction' would both reduce and herald leakage, leading to a so-called erasure error [100, 101]. Unlike Pauli errors, the exact location of erasures is known, making them easier to correct and leading to higher error thresholds [102–105].

Here, we present the realization and extension of the LRU scheme proposed in Ref. [72]. This is a highly practical scheme requiring only microwave pulses and the quantum hardware typically found in contemporary circuit QED quantum processors: a microwave drive and a readout resonator dispersively coupled to the target transmon (in our case, a readout res-

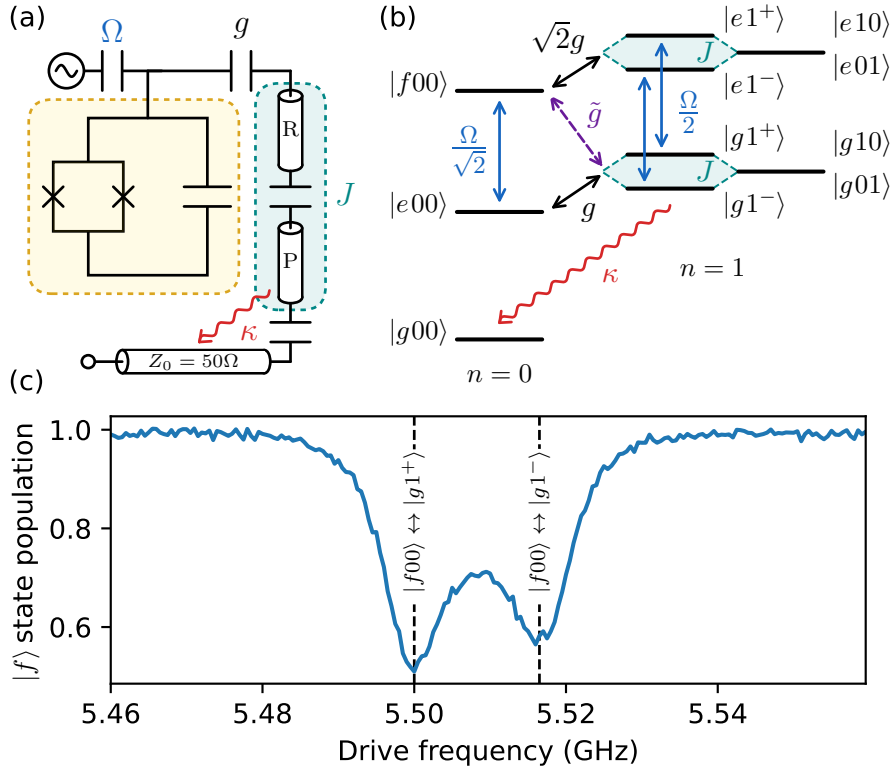


Figure 4.1: **Leakage reduction unit scheme.** (a) Schematic for the driven transmon-resonator system. A transmon (T, yellow) with three lowest-energy levels $|g\rangle$, $|e\rangle$, and $|f\rangle$ is coupled to a readout resonator (R) with strength g . The latter is coupled to a frequency-matched Purcell resonator (P) with strength J . The Purcell resonator also couples to a $50\ \Omega$ feedline through which its excitations quickly decay at rate κ . The transmon is driven with a pulse of strength Ω applied to its microwave drive line. (b) Energy level-spectrum of the system. Levels are denoted as $|T, R, P\rangle$, with numbers indicating photons in R and P. As the two resonators are frequency matched, the right-most degenerate states split by $2J$, and g is shared equally among the two hybridized resonator modes $|1^-\rangle$ and $|1^+\rangle$. An effective coupling $\tilde{g}$ arises between $|f00\rangle$ and the two hybridized states $|g1^\pm\rangle$ via $|e00\rangle$ and $|e1^\pm\rangle$. (c) Spectroscopy of the $|f00\rangle \leftrightarrow |g1^\pm\rangle$ transition. Measured transmon population in $|f\rangle$ versus drive frequency, showing dips corresponding to the two transitions assisted by each of the hybridized resonator modes.

onator with dedicated Purcell filter). We show its straightforward calibration and the effective removal of the population in the first two leakage states of the transmon ($|f\rangle$ and $|h\rangle$) with up to $> 99\%$ efficacy in 220 ns. Process tomography reveals that the LRU backaction on the qubit subspace is only an AC-Stark shift, which can be easily corrected using a Z -axis rotation. As a first application in a QEC setting, we interleave repeated measurements of a weight-2 parity check [57, 58] with simultaneous LRUs on data and ancilla qubits, showing the suppression of leakage and error detection rate buildup.

4.2 Results

4.2.1 All-microwave leakage reduction unit scheme

Our leakage reduction scheme [Fig. 4.1(a)] consists of a transmon with states $|g\rangle$, $|e\rangle$ and $|f\rangle$, driven by an external drive Ω , coupled to a resonant pair of Purcell and readout resonators [41] with effective dressed states $|00\rangle$ and $|1^\pm\rangle$. The LRU scheme transfers leakage population in the second-excited state of the transmon, $|f\rangle$, to the ground state, $|g\rangle$, via the resonators using a microwave drive. It does so using an effective coupling, $\tilde{g}$, mediated by the transmon-resonator coupling, g , and the drive Ω , which couples states $|f00\rangle$ and $|g1^\pm\rangle$. Driving at the frequency of this transition,

$$\omega_{f00} - \omega_{g1^\pm} \approx 2\omega_Q + \alpha - \omega_{RP}, \quad (4.1)$$

transfers population from $|f00\rangle$ to $|g1^\pm\rangle$, which in turn quickly decays to $|g00\rangle$ provided the transition rate, $\tilde{g}$, is small compared to κ . Here, ω_Q and α are the transmon qubit transition frequency and anharmonicity, respectively, while ω_{RP} is the resonator mode frequency. In this regime, the drive effectively pumps any leakage in $|f\rangle$ to the computational state $|g\rangle$. We perform spectroscopy of this transition by initializing the transmon in $|f\rangle$ and sweeping the drive around the expected frequency. The results [Fig. 4.1(c)] show two dips in the f -state population corresponding to transitions with the hybridized modes of the matched readout-Purcell resonator pair. The dips are broadened by $\sim \kappa_{\text{eff}}/2\pi \approx 8$ MHz, where $\kappa_{\text{eff}} = \kappa/2$ is the effective linewidth of the dressed resonator (see section 4.4 for device characteristics and metrics), making them easy to find. We achieve typical couplings of $\tilde{g}/2\pi \sim 1$ MHz for this transition (sec. 4.4).

4.2.2 Calibration of the leakage reduction unit pulse

To make use of this scheme for a LRU, we calibrate a pulse that can be used as a circuit-level operation. We use the pulse envelope proposed in Ref. [72]:

$$A(t) = \begin{cases} A \sin^2\left(\pi \frac{t}{2t_r}\right) & \text{for } 0 \leq t \leq t_r, \\ A & \text{for } t_r \leq t \leq t_p - t_r, \\ A \sin^2\left(\pi \frac{t_p - t}{2t_r}\right) & \text{for } t_p - t_r \leq t \leq t_p, \end{cases} \quad (4.2)$$

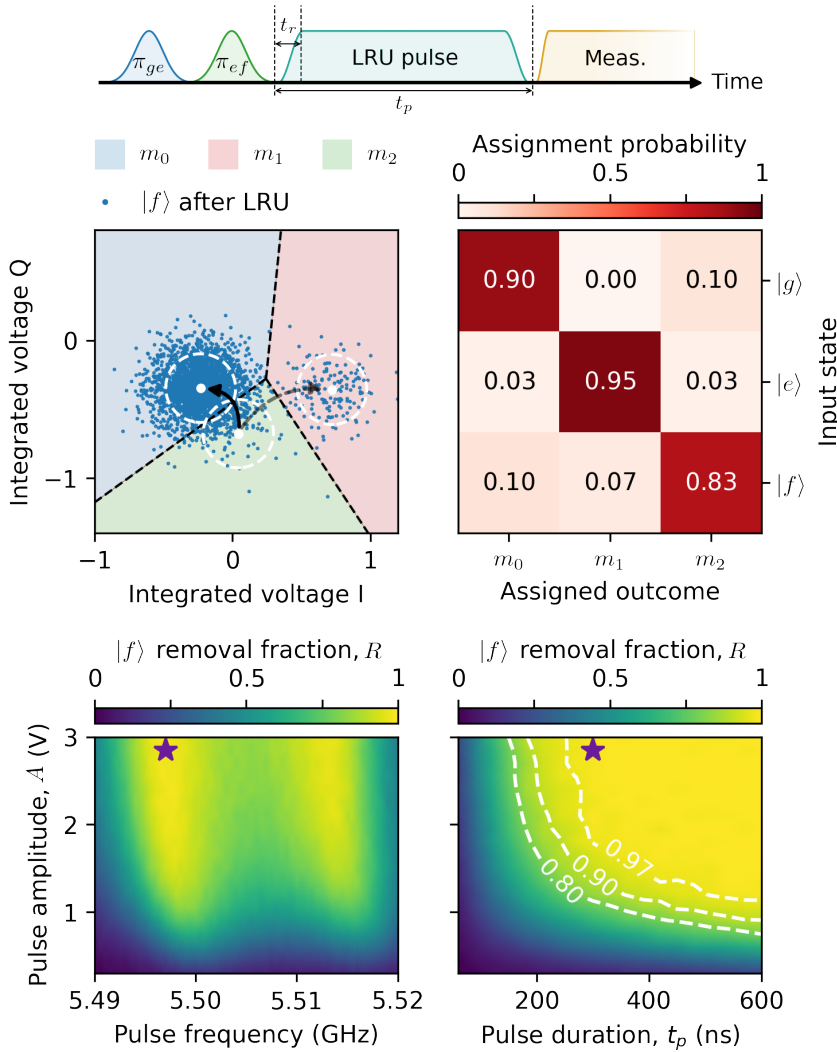


Figure 4.2: **Calibration of the leakage reduction unit pulse.** (a) Pulse sequence used for LRU calibration. (b) Single-shot readout data obtained from the experiment. The blue, red and green areas denote m_0 , m_1 , and m_2 assignment regions, respectively. The mean (white dots) and 3σ standard deviation (white dashed circles) shown are obtained from Gaussian fits to the three input-state distributions. The blue data shows the first 3×10^3 (from a total of 2^{15}) shots of the experiment described in (a), indicating 99.(3)% $|f\rangle$ -state removal fraction. (c) Measured assignment fidelity matrix used for readout correction. (d-e) Extracted $|f\rangle$ -state removal fraction versus pulse parameters. Added contours (white dashed curves) indicate 80, 90 and 97% removal fraction. The purple star indicates the pulse parameters used in (b).

where A is the amplitude, t_r is the rise and fall time, and t_p is the total duration. We conservatively choose $t_r = 30$ ns to avoid unwanted transitions in the transmon. To measure the fraction of leakage removed, R , we apply the pulse on the transmon prepared in $|f\rangle$ and measure it [Fig. 4.2(a)], correcting for readout error using the measured 3-level assignment fidelity matrix [Fig. 4.2(c)]. To optimize the pulse parameters, we first measure R while sweeping the pulse frequency and A [Fig. 4.2(d)]. A second sweep of t_p and A [Fig. 4.2(e)] shows that $R > 99\%$ can be achieved by increasing either parameters. This value is limited by thermal population in the resonator modes. We estimate values of $P(n = 1) \approx 0.5\%$ (sec. 4.4). Simulation [72] suggests that $R \approx 80\%$ is already sufficient to suppress most of the impact of current leakage rates, which is comfortably achieved over a large region of parameter space. For QEC, a fast operation is desirable to minimize the impact of decoherence. However, one must not excessively drive the transmon, which can cause extra decoherence (see Fig. 6 in Ref. [72]). Considering the factors above, we opt for $t_p = 220$ ns and adjust A such that $R \gtrsim 80\%$. Additionally, we benchmark the repeated action of the LRU and verify that its performance is maintained over repeated applications, thus restricting leakage events to approximately a single cycle (see section 4.4 Fig. S2).

4.2.3 Benchmarking in the qubit subspace

With the LRU calibrated, we then benchmark its impact on the qubit subspace using quantum process tomography. The results (Fig. 4.3) show that the qubit incurs a Z -axis rotation. We find that the rotation angle increases linearly with t_p [Fig. 4.3(g)], consistent with a 71(9) kHz AC-Stark shift induced by the LRU drive. This phase error in the qubit subspace can be avoided using decoupling pulses or corrected with a virtual Z gate. Figures 4.3(h) and 4.3(i) show the Pauli transfer matrix (PTM) for the operation before and after applying a virtual Z correction, respectively. From the measured PTM [Fig. 4.3(i)] and enforcing physicality constraints [64], we obtain an average gate fidelity $F_{\text{avg}} = 98.(9)\%$. Compared to the measured 99.(2)% fidelity of idling during the same time ($t_p = 220$ ns), there is evidently no significant error increase.

4.2.4 Reduction of leakage in repeated stabilizer measurements

Finally, we implement the LRU in a QEC scenario by performing repeated stabilizer measurements of a weight-2 X -type parity check [57, 58] using three transmons (Fig. 4.4). We use the transmon in Figs. 1-3, D_1 , plus an additional transmon (D_2) as data-qubits together with an ancilla, A . LRUs for D_2 and A are tuned using the same procedure as above. A detailed study of the performance of this parity check and of the impact of simultaneous LRUs is shown in section 4.4 Figs. S6 and S5. Given their frequency configuration [21], D_1 and A are most vulnerable to leakage during two-qubit controlled- Z (CZ) gates, as shown by the avoided crossings in Fig. 4.4(a). Additional leakage can occur during other operations: in particular, we observe that leakage into states above $|f\rangle$ can occur in A due to measurement-induced

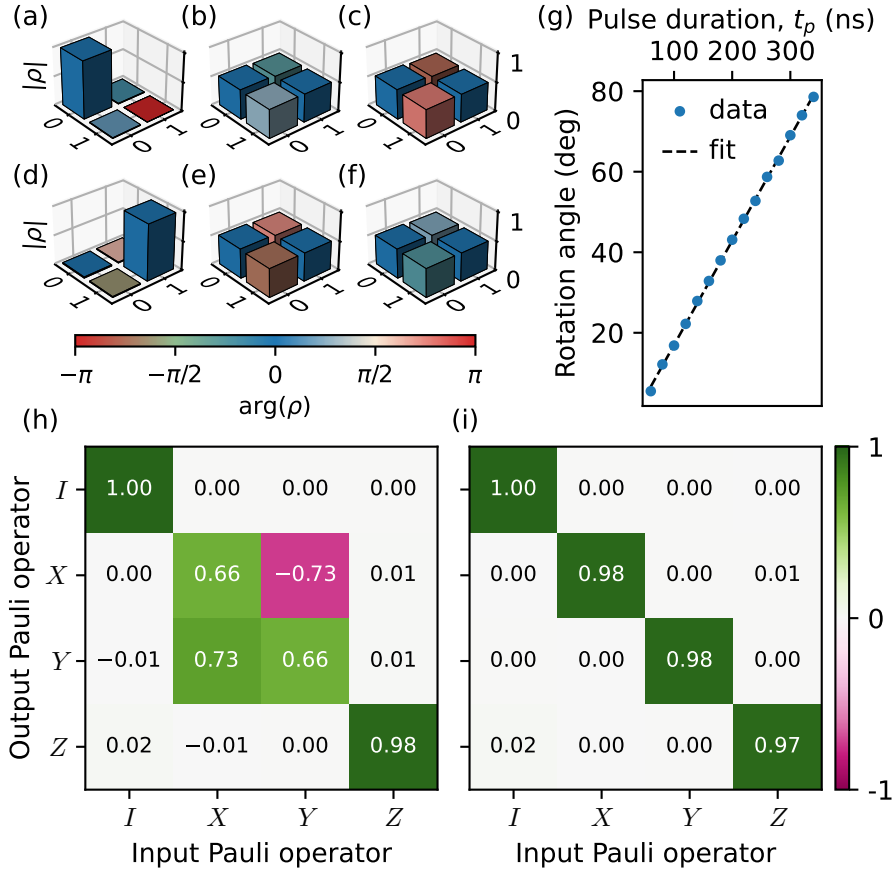


Figure 4.3: **Process tomography of the leakage reduction unit.** (a-f) Measured density matrices after the LRU gate for input states $|0\rangle$, $|+\rangle$, $|+i\rangle$, $|1\rangle$, $|-\rangle$, and $| -i\rangle$, respectively. (g) Z-rotation angle induced on the qubit versus the LRU pulse duration. The linear best fit (black dashed line) indicates an AC-Stark shift of 71(9) kHz. (h-i) Pauli transfer matrix of the LRU with (i) and without (h) virtual phase correction ($t_p = 220$ ns and $R = 84.(7)\%$).

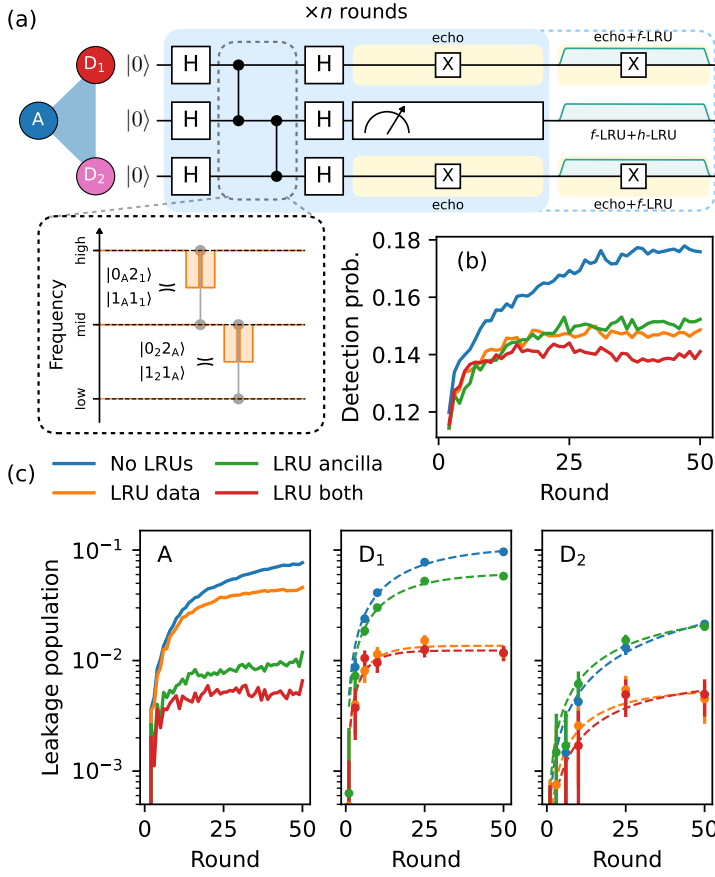


Figure 4.4: **Repeated stabilizer measurement with leakage reduction.** (a) Quantum circuit using ancilla A to measure the X-type parity of data qubits D1 and D2. The dashed box shows the frequency arrangement for two-qubit CZ gates. A CZ gate is performed by fluxing the higher-frequency transmon down in frequency to the nearest avoided crossing (orange shaded trajectory). The duration of single-qubit gates, CZ gates, and measurement are 20, 60 and 340 ns, respectively, totalling 500 ns for the parity check (light-blue region). Performing the LRUs extends the circuit by 220 ns (blue-dashed region). Echo pulses on data qubits mitigate phase errors caused by residual ZZ crosstalk and AC-stark shift during the measurement and LRUs (light yellow slots). (b-c) Measured error-detection probability (b) and leakage (c) versus the number of parity-check rounds in four settings. Here, leakage includes any population outside the computational subspace. The No LRUs setting (blue) does not apply any LRUs. LRU data (orange) and LRU ancilla (green) settings apply LRUs exclusively on the data qubits and the ancilla, respectively. The LRU both (red) setting applies LRUs on all qubits.

transitions [18] (see section 4.4 Fig. S10). Therefore, a LRU acting on $|f\rangle$ alone is insufficient for A. To address this, we develop an additional LRU for $|h\rangle$ (h -LRU), the third-excited state of A (see section 4.4 Fig. S9). The h -LRU can be employed simultaneously with the f -LRU without additional cost in time or impact on the $|f\rangle$ removal fraction, R . Thus, we simultaneously employ f -LRUs for all three qubits and an h -LRU for A [Fig. 4.4(a)]. To evaluate the impact of the LRUs, we measure the error detection probability (probability of a flip occurring in the measured stabilizer parity) and leakage population of the three transmons over multiple rounds of stabilizer measurement. Without leakage reduction, the error detection probability rises $\sim 8\%$ in 50 rounds. We attribute this feature to leakage build-up [39, 71, 99]. With the LRUs, the rise stabilizes faster (in ~ 10 rounds) to a lower value and is limited to 2% , despite the longer cycle duration (500 versus 720 ns without and with the LRU, respectively). Leakage is overall higher without LRUs, in particular for D_1 and A [Fig. 4.4(c)], which show a steady-state population of $\approx 10\%$. Using leakage reduction, we lower the leakage steady-state population by up to one order of magnitude to $\lesssim 1\%$ for all transmons. Additionally, we find that removing leakage on other transmons leads to lower overall leakage, suggesting that leakage is transferred between transmons [70, 99]. This is particularly noticeable in A [Fig. 4.4(c)], where the steady-state leakage is always reduced by adding LRUs on D_1 and D_2 .

4.3 Discussion

We have demonstrated and extended the all-microwave LRU for superconducting qubits in circuit QED proposed in Ref. [72]. We have shown how these LRUs can be calibrated using a straightforward procedure to deplete leakage in the second- and third-excited states of the transmon. This scheme could potentially work for even higher transmon states using additional drives. We have verified that the LRU operation has minimal impact in the qubit subspace, provided one can correct for the static AC-Stark shift induced by the drive(s).

This scheme does not reset the qubit state and is therefore compatible with both data and ancilla qubits in the QEC context. We have showcased the benefit of the LRU in a building-block QEC experiment where LRUs decrease the steady-state leakage population of data and ancilla qubits by up to one order of magnitude (to $\lesssim 1\%$), and thereby reduce the error detection probability of the stabilizer and reaching a faster steady state. We find that the remaining ancilla leakage is dominated by higher states above $|f\rangle$ (see section 4.4 Fig. S10) likely caused by the readout [18, 19]. Given the observation leakage transfer between transmons, which can result in higher excited leakage states [99], data qubits can also potentially benefit from h -LRUs. Compared to other LRU approaches [71, 99], we believe this scheme is especially practical as it is all-microwave and very quantum-hardware efficient, requiring only the microwave drive line and dispersively coupled resonator that are already commonly found in the majority of circuit QED quantum processors [38, 39, 89]. Extending this leakage reduction method to larger QEC experiments can be done without further penalty in time as

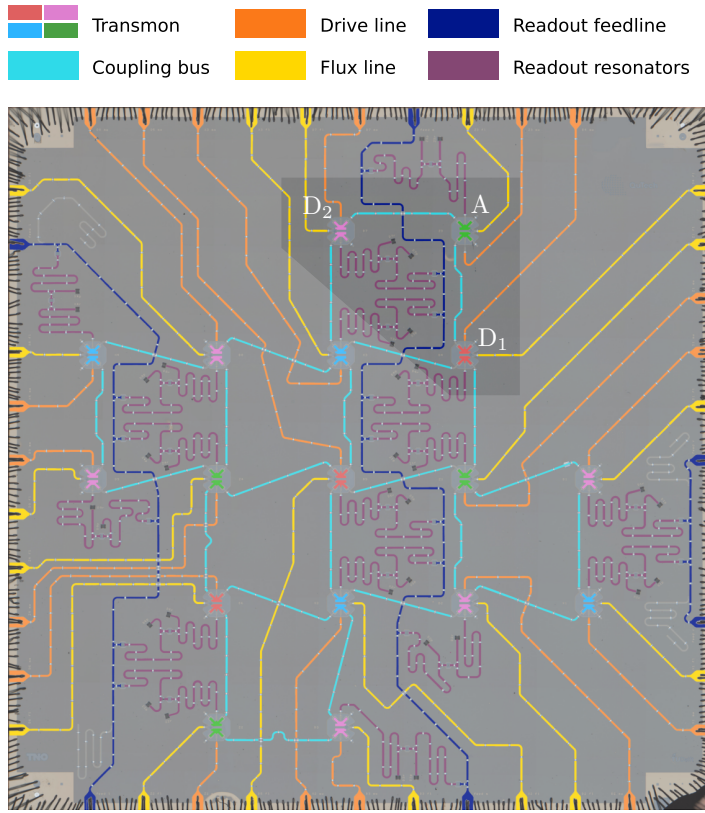


Figure 4.5: **Circuit QED device.** Optical image of the 17-transmon quantum processor, with added falsecolor to highlight different circuit elements. The shaded area indicates the three transmons used in this experiment.

all LRUs can be simultaneously applied. However, we note that when extending the LRU to many qubits, microwave crosstalk should be taken into account in order to avoid driving unwanted transitions. This can be easily avoided by choosing an appropriate resonator-qubit detuning.

4.4 Supplementary information

This supplement provides additional information in support of the statements and claims in this chapter.

Transmon	D ₁	A	D ₂
Frequency at sweetspot, $\omega_Q/2\pi$ (GHz)	6.802	6.033	4.788
Anharmonicity, $\alpha/2\pi$ (MHz)	-295	-310	-321
Resonator frequency, $\omega_{R/P}/2\pi$ (GHz)	7.786	7.600	7.105
Purcell res. linewidth, $\kappa/2\pi$ (MHz)	15.5	22.5	12.6
Qubit-res. coupling, $g/2\pi$ (MHz)	172	212	301
f -LRU drive frequency (GHz)	5.498	4.135	2.152
h -LRU drive frequency (GHz)	-	3.496	-
T_1 (τ s)	17	26	37
$T_{2,\text{echo}}$ (τ s)	19	22	27
Single-qubit gate error (%)	0.1(0)	0.0(7)	0.0(5)
Two-qubit gate error (%)	1.(1)	1.(9)	
Two-qubit gate leakage (%)	0.3(7)	0.1(1)	
f -LRU removal fraction, R^f (%)	84.(7)	99.(2)	80.(3)
h -LRU removal fraction, R^h (%)	-	96.(1)	-
Operation	Duration (ns)		
Single-qubit gate	20		
Two-qubit gate	60		
Measurement	340		
LRU	220		

Table 4.1: **Device metrics.** Frequencies and coherence times are measured using standard spectroscopy and time-domain measurements [13]. Gate errors are evaluated using randomized benchmarking protocols [26, 27, 81].

4.4.1 Device

The device used (Fig. 4.5) has 17 flux-tunable transmons arranged in a square lattice with nearest-neighbor connectivity (as required for a distance-3 surface code). Transmons are arranged in three frequency groups as prescribed in the pipelined architecture of Ref. [21]. Each transmon has a dedicated microwave drive line used for single-qubit gates and leakage reduction, and a flux line used for two-qubit gates. Nearest-neighbor transmons have fixed coupling mediated by a dispersively coupled bus resonator. Each transmon has a dedicated pair of frequency-matched readout and Purcell resonators coupled to one of three feedlines, used for fast multiplexed readout in the architecture of Ref. [41]. Single-qubit gates are realized using standard DRAG pulses [23]. Two-qubit controlled- Z gates are implemented using sudden net-zero flux pulses [11]. Characteristics and performance metrics of the three transmons used in the experiment are shown in Tab. 4.1.

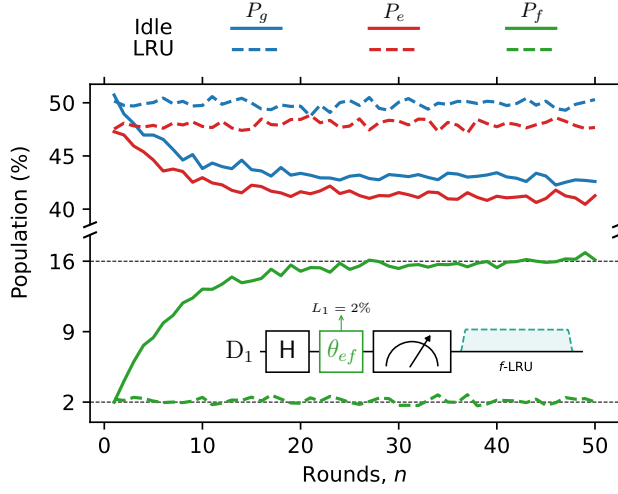


Figure 4.6: **Repeated LRUs on D_1** . Computational (top) and leakage (bottom) state population over repeated cycles of the circuit as shown. Transmon D_1 is repeatedly put in superposition, controllably leaked with rate L_1 and measured. Measurement is followed by either idling (solid) or the LRU (dashed). The horizontal lines denote the steady-state leakage with and without LRUs.

4.4.2 Repeated LRU application

For QEC we require that the LRU performance remains constant over repeated applications. To assess this, we perform repeated rounds of the experiment shown in Fig. 4.6 while idling or using the LRU. In each round apply an e - f rotation with rotation angle θ to induce a leakage rate

$$L_1 = \frac{\sin^2(\theta/2)}{2}, \quad (4.3)$$

and choose θ such that $L_1 = 2\%$. For the purpose of this experiment, we lower the readout amplitude in order to suppress leakage to higher states during measurement [18]. The results (Fig. 4.6) show that while idling, leakage in $|f\rangle$ builds up to a steady-state population of about 16%. Using the LRU, it remains constant at $P_f = L_1$ throughout all rounds. This behavior shows that LRU performance is maintained throughout repeated applications and suggests that leakage events are restricted to a single round.

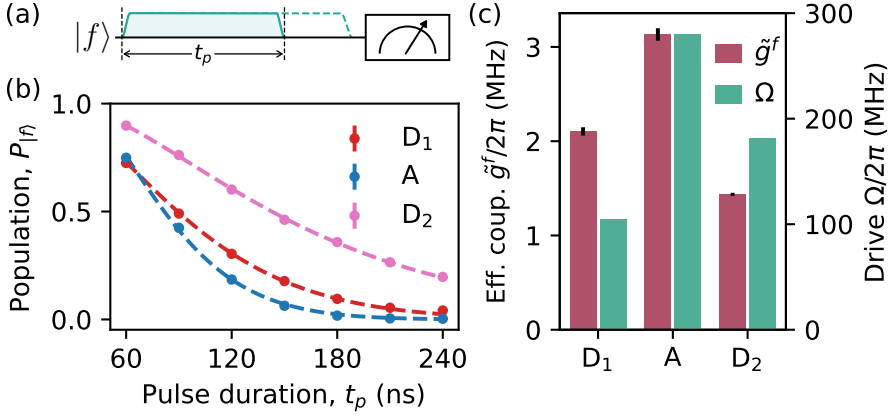


Figure 4.7: **Assessing the $|f00\rangle \leftrightarrow |g1^\pm\rangle$ transition coupling.** (a) Experimental circuit to measure P_f versus LRU pulse duration, t_p . After initializing the transmon in $|f\rangle$, we perform a LRU pulse of duration t_p within a constant interval of 240 ns followed by a measurement. (b) Measured $|f\rangle$ state population, P_f , as a function of t_p for each of the three transmons used. The dashed lines show fits to the data used to estimate the effective transition coupling $\tilde{g}^f$. (c) Transition couplings, $\tilde{g}^f$, and corresponding drive strengths, Ω , estimated from fitting.

4.4.3 Estimating effective transition coupling $\tilde{g}^f$

In order to assess the effective transition coupling, $\tilde{g}^f$, we study the system described in Fig. 1(a). Excluding the drive, the total Hamiltonian of this system is given by

$$H = H_R + H_P + H_Q + J(a_R^\dagger a_P + a_R a_P^\dagger) + g(a_R^\dagger b + a_R b^\dagger), \quad (4.4)$$

where $H_{R/P} = \omega_{R/P} a^\dagger a$ and $H_Q = \omega_Q b^\dagger b + \frac{\alpha}{2} b^\dagger b^\dagger b b$ are the resonators and transmon Hamiltonians, respectively. When the two resonator modes are resonant, i.e., $\omega_R = \omega_P$,

$$H = H_R^+ + H_R^- + H_Q + \frac{g}{\sqrt{2}}[(a_+^\dagger b + a_+ b^\dagger) + (a_-^\dagger b + a_- b^\dagger)], \quad (4.5)$$

where $a_\pm = (a_R \pm a_P)/\sqrt{2}$ and $H_R^\pm = (\omega_R \pm J)a_\pm^\dagger a_\pm$ describe the dressed resonator modes. In this basis, we find that g is shared by the dressed resonator modes with effective strength $g_{\text{eff}} = g/\sqrt{2}$. Because of this, each dressed resonator mode approximately inherits an effective linewidth $\kappa_{\text{eff}} = \kappa/2$. We then study the dynamics of this system in the three-level manifold $\{|g00\rangle, |g1^\pm\rangle, |f00\rangle\}$ by solving the Lindblad equation

$$\dot{\rho} = -i[H, \rho] + L\rho L^\dagger - \frac{1}{2}\{L^\dagger L, \rho\}, \quad (4.6)$$

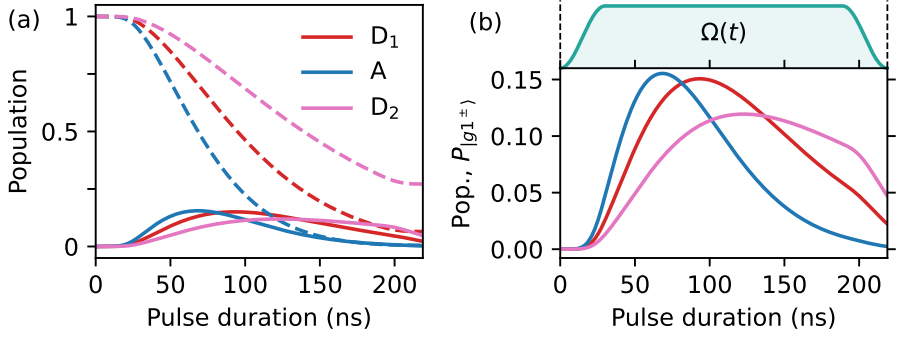


Figure 4.8: **Dynamics of the states $|f00\rangle$ and $|g1^\pm\rangle$ during the LRU.** (a) Population of the state $|f00\rangle$ (dashed) and $|g1^\pm\rangle$ (solid) during the LRU pulse, $\Omega(t)$, obtained by numerically solving Eq. 4.6. (b) Close-up of population in $|g1^\pm\rangle$ and LRU pulse (top) versus time.

where $L = \sqrt{\kappa_{\text{eff}}} |g00\rangle \langle g1^\pm|$ is the decay operator modeling loss in the resonator mode and

$$H = \begin{bmatrix} 0 & 0 & 0 \\ 0 & 0 & \tilde{g}^f \\ 0 & (\tilde{g}^f)^* & 0 \end{bmatrix} \quad (4.7)$$

is the approximate Hamiltonian of the system in the drive frame on resonance with $|f00\rangle \leftrightarrow |g1^\pm\rangle$. In this model, an initial state $|f00\rangle$ evolves as (Eq. S2 of Ref. [92]),

$$P_{|f00\rangle}(t) = e^{-\frac{\kappa_{\text{eff}}}{2}t} \left| \cosh\left(\frac{\zeta}{2}t\right) + \frac{\kappa_{\text{eff}}}{2\zeta} \sinh\left(\frac{\zeta}{2}t\right) \right|^2, \quad (4.8)$$

where, $\zeta = \sqrt{-(2\tilde{g}^f)^2 + (\kappa_{\text{eff}}/2)^2}$. To estimate $\tilde{g}^f$, we measure $P_{|f\rangle}$ as function of pulse duration, t_p , as shown in Fig. 4.7 and fit

$$P_{|f\rangle}(t_p, \kappa_{\text{eff}}, \tilde{g}^f, t_0) = P_{|f00\rangle}(t_p - t_0), \quad (4.9)$$

where t_0 is introduced to account for the finite rise and fall time of the LRU drive. We fix $\kappa_{\text{eff}} \equiv \kappa/2$ (shown in Tab. 4.1) leaving only two parameters to fit $\tilde{g}^f$, and t_0 . The results [Fig. 4.7(c)] show $\tilde{g}^f/2\pi$ lies between 1 and 3 MHz. Additionally, we find $t_0 \approx 20$ ns similar to the rise time, $t_r = 30$ ns, used in experiment. From these results, one can also estimate the drive strength Ω used for each transmon using (Eq. A35 of Ref. [72]),

$$\tilde{g}^f = \Omega \frac{g_{\text{eff}}\alpha}{\sqrt{2\Delta(\Delta + \alpha)}}, \quad (4.10)$$

where $\Delta \equiv \omega_Q - \omega_R$ and $g_{\text{eff}} \equiv g/\sqrt{2}$. We find drive strengths, $\Omega/2\pi$, lying between 100 and 300 MHz depending on the qubit. This disparity occurs because the drives are synthesized at distinct frequencies (from 2 to 5.5 GHz) and therefore use different RF components

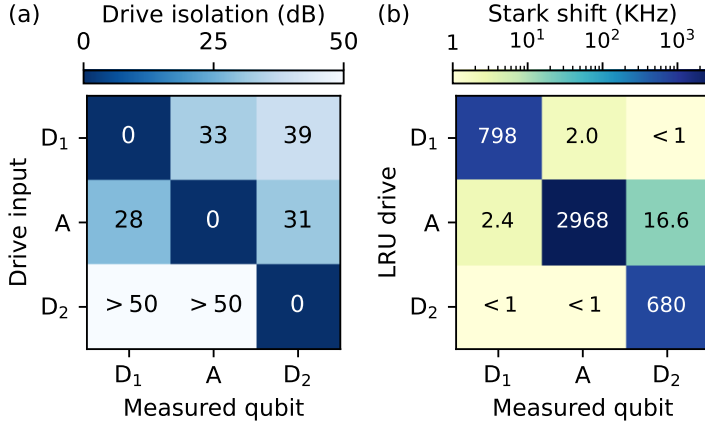


Figure 4.9: **Cross-driving induced AC Stark shift from LRU drives.** (a) Measured drive isolation between the qubits. For each element in the matrix S_{ij} , we measure the isolation of qubit i to a drive Ω_j driven through the drive line of qubit j such that its effective strength is $\Omega_i^{\text{eff}} = 10^{-S_{ij}/20} \Omega_j$. (b) Estimated AC Stark shift induced by the LRU drive between qubits based on the drive amplitudes calculated in Fig. 4.7(c). We plot absolute frequency shifts, however, actual estimated values are negative.

and experience different attenuation. We note that the amplitudes used in Fig. 4.7 are slightly different from the ones used in the experiments reported in the main text however they lead to similar removal fractions.

4.4.4 Population in the readout resonators

In this section, we consider the impact of population in the resonator modes. Thermal population in the resonators can converted into leakage during the LRU [72]. As shown in Fig. 2(d) of Ref. [72], when there is no leakage, the system will evolve such that $P_{|f00\rangle} = P_{|g1^\pm\rangle}$ in the steady state. It is thus important that the mean photon number, $\bar{n}$, characterising the thermal field in the resonators is low, $\bar{n} \ll 1$. One can estimate an upper bound of $\bar{n}$ assuming that the pure dephasing rate of the transmon, Γ_ϕ , is limited by photon shot noise in the readout resonators using [13], namely

$$\Gamma_\phi = \eta \frac{4\chi^2}{\kappa} \bar{n}, \quad (4.11)$$

where $\eta = \kappa^2 / (\kappa^2 + 4\chi^2)$. Using this estimate, we find $P(n=1) \approx 0.5\%$. Therefore, one can expect at most $P_{|f00\rangle} \approx 0.5\%$ in the steady state, which will limit the removal fraction $R < 99.5\%$. Another important aspect to consider when using the LRU repeatedly, is the amount of population left in $|g1^\pm\rangle$ after a leakage event. If significant population is left in

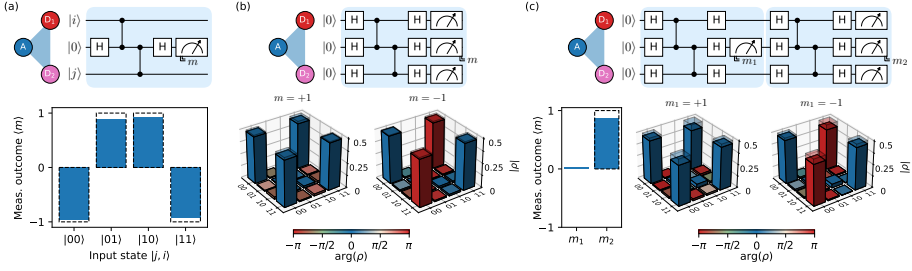


Figure 4.10: Benchmarking of the weight-2 parity check. (a) Quantum circuit of the weight-2 X-type parity check and bar plot of the average measured ancilla outcome for the different input computational states of the data-qubit register. Dashed bars show ideal average outcome: $\langle m \rangle = +1$ (-1) for even (odd) data-qubit input parity. (b) Generation of Bell states via stabilizer measurement (top) and corresponding data-qubit state tomography (bottom) conditioned on the ancilla outcome. The obtained fidelity to the ideal Bell states (shaded wireframe) is 96.9% and 98.5% for $m = +1$ and $m = -1$, respectively. (c) Repeated stabilizer experiment. (Bottom left) Average measured ancilla outcome $\langle m \rangle$ for each round of stabilizer measurement. Ideally, the first outcome should be random and the second always $+1$. The measured probability is $P(m_2 = +1) = 90.0\%$. (Bottom right) Reconstructed data-qubit states conditioned on the first ancilla outcome. The obtained Bell-state fidelities are 90.6% and 91.5% for $m_1 = +1$ and $m_1 = -1$, respectively.

the resonator, errors can occur in subsequent operations. We investigate this by numerically solving Eq. 4.6 taking into account the rise time of the LRU pulse and using the $\tilde{g}$ estimated in Fig. 4.7. Figure 4.8 shows the evolution of state $|g1^\pm\rangle$ for the three transmons. We find that the population in this level never exceeds 15%. Moreover, Fig. 4.8(b) reveals that the leftover population by the end of the pulse is below 5%. In particular, we see that population is significantly suppressed during the pulse fall time t_f . Therefore, we do not expect errors following the LRU to be significant after a leakage event.

4.4.5 Drive crosstalk

Given the relatively high strength of the drives required to drive the LRUs, we consider the impact of drive crosstalk. In particular, we study the impact of each LRU drive on each qubit. By coupling off-resonantly, cross-drives may shift the qubit frequency which could lead to lower performance of the LRUs when driven simultaneously. To investigate this, we first characterize drive crosstalk between the three transmons used by measuring the drive isolation,

$$S_{ij} = 20 \log_{10}(\Omega_j/\Omega_i^{\text{eff}}), \quad (4.12)$$

between a drive, Ω_j , on qubit j and the effective drive, Ω_i^{eff} , felt on qubit i . The results [Fig. 4.9(a)] show that drive isolation is at least 28 dB and crosstalk is most significant when

driving via the drive lines of D_1 and A. Using the drive strength estimated in Fig. 4.7, we then calculate the AC Stark shift, δ , induced by an effective drive Ω_{eff} using,

$$\delta = \frac{\alpha \Omega_{\text{eff}}^2}{2\Delta_d(\Delta_d + \alpha)}, \quad (4.13)$$

where $\Delta_d = \omega_{\text{drive}} - \omega_Q$. The results [Fig. 4.9(b)] show that crosstalk induced shifts are at most ≈ 17 kHz. This is much smaller than the typical linewidth of the $|f00\rangle \leftrightarrow |g1^\pm\rangle$ transition ~ 10 MHz [as shown in Fig. 1(c)]. Therefore, one does not expect lower simultaneous performance of LRUs due to cross-drive Stark shifts. One can also verify this experimentally by measuring simultaneous leakage removal fractions, R , and comparing them to individual ones. Performing this experiment, we find the same removal fractions, as expected.

4.4.6 Benchmarking the weight-2 parity check

We benchmark the performance of the weight-2 parity check using three experiments assessing different error types. First, we assess the ability to accurately assign the parity of the data-qubit register by measuring the ancilla outcome for all data-qubit input computational states. The results [Fig. 4.10(a)] show an average parity assignment fidelity of 95.6%. Next, we look at errors occurring on the data qubits when projecting them onto a Bell state using a X -type parity check [Fig. 4.10(b)]. From data-qubit state tomography conditioned on ancilla outcome, we obtain an average Bell-state fidelity of 97.7% (96.9% for $m = +1$ and 98.5% for $m = -1$). For each reported density matrix, we apply readout corrections and enforce physicality constraints via maximum likelihood estimation [64]. Finally, we look at the backaction of two back-to-back parity checks [Fig. 4.10(c)]. Here, we measure the correlation of the two parity outcomes. Ideally, the first parity outcome should be random while the second should be the same as the first. Since our ancilla is not reset after measurement, the probability of both parities being correlated is $P(m_2 = +1) = 90.0\%$ [bar plot in Fig. 4.10(c)]. We can also reconstruct the data qubit state after the experiment. Here, we find that the average Bell-state fidelity drops to 91.0% (90.6% for $m_1 = +1$ and 91.5% for $m_1 = -1$). This drop in fidelity is likely due to decoherence from idling during the first ancilla measurement.

4.4.7 Measurement-induced transitions

Previous studies have observed measurement-induced state transitions that can lead to leakage [18, 19]. To evaluate the backaction of ancilla measurement, we model the measurement as a rank 3 tensor $\epsilon_i^{m,j}$ which takes an input state i , declares an outcome m and outputs a state j with normalization condition,

$$\sum_m \sum_j \epsilon_i^{m,j} = 1. \quad (4.14)$$

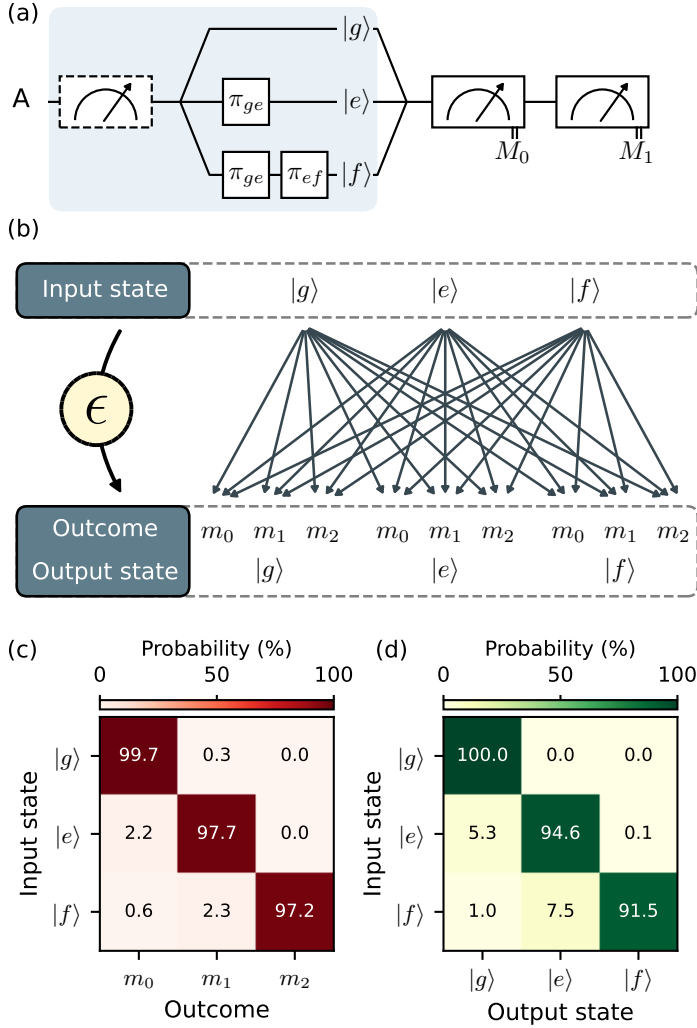


Figure 4.11: **Characterizing measurement-induced transitions.** (a) Quantum circuit used to characterize transmon measurement. A transmon is initialized into states $|g\rangle$, $|e\rangle$ and $|f\rangle$ after a heralding (dashed) measurement (blue panel). Following preparation, two consecutive measurements M_0 and M_1 are performed, yielding three-level outcomes. (b) Illustration of the extracted measurement model. The model is described by a rank 3 tensor ϵ_i^{mj} , where input states i are connected to measurement outcomes m and output states j . From it, the assignment probability matrix (c) and the QNDness matrix (d) can be extracted.

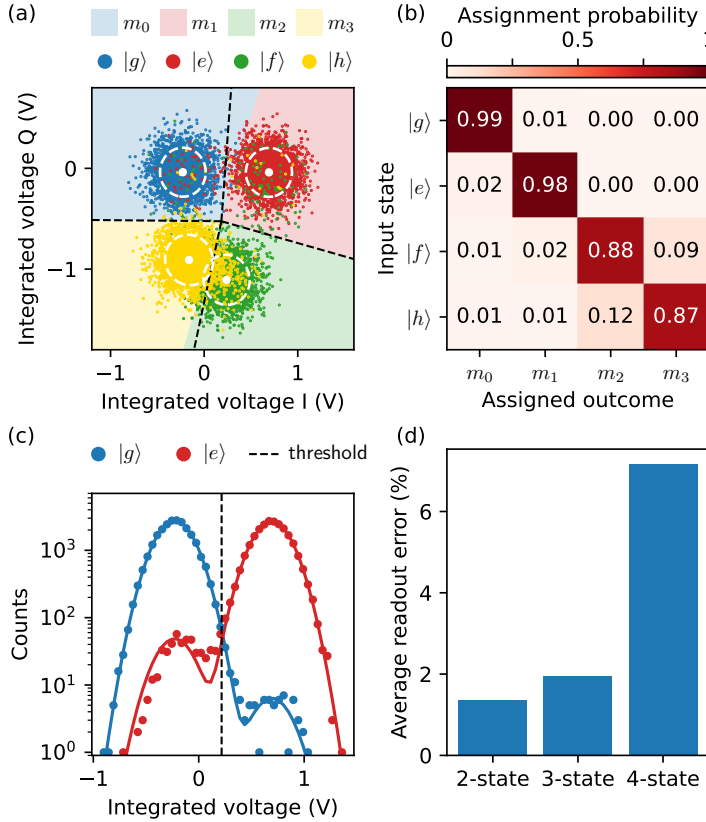


Figure 4.12: **Four state readout.** (a) Single-shot readout data for the four lowest-energy transmon states $|g\rangle$, $|e\rangle$, $|f\rangle$ and $|h\rangle$ of A. Data are plotted for the first 3×10^3 from a total of 2^{15} shots for each input state. The dashed lines show decision boundaries obtained from fitting a linear discrimination classifier to the data. The mean (white dot) and 3σ standard deviation (white dashed circles) shown are obtained from Gaussian fits to each input state distribution. (b) Assignment probability matrix obtained from classification of each state into a quaternary outcome. (c) Histogram of shots for qubit states taken along the projection maximizing the signal-to-noise ratio. (d) The average assignment errors for 2-, 3- and 4-state readout are 1.(3), 1.(9) and 7.(2)%, respectively.

To find ϵ_i^{mj} , we perform the experiment in Fig. 4.11(a). For each input state i , the probability distribution of the measured results $P_i(M_0, M_1)$ follows

$$P_i(M_0 = m_k, M_1 = m_\ell) = \sum_s \sum_j \epsilon_i^{m_k, s} \epsilon_s^{m_\ell, j}. \quad (4.15)$$

This system of 27 second-order equations is used to estimate the 27 elements of ϵ_i^{mj} through a standard optimization procedure. In this description, the assignment fidelity matrix $M_{i,m}$ [Fig. 4.11(b)] is given by

$$M_{i,m} = \sum_j \epsilon_i^{mj}. \quad (4.16)$$

Furthermore, this model allows us to assess the probability of transitions occurring during the measurement. This is given by the QNDness matrix

$$Q_{i,j} = \sum_m \epsilon_i^{mj}. \quad (4.17)$$

The results [Fig. 4.11(b)] show an average QNDness of 95.4% across all states. The average leakage rate $((Q_{g,f} + Q_{e,f})/2)$ is 0.06%, predominantly occurring for input state $|e\rangle$.

4.4.8 Readout of transmon states

In order to investigate leakage to higher states in the ancilla, we need to discriminate between the first two leakage states, $|f\rangle$ and $|h\rangle$. To do this without compromising the performance of the parity check, we simultaneously require high readout fidelity for the qubit states $|g\rangle$ and $|e\rangle$. We achieve this for the ancilla for the states $|g\rangle$ through $|h\rangle$ using a single readout pulse. Figure 4.12(a) shows the integrated readout signal for each of the states along with the decision boundaries used to classify the states. Any leakage to even higher states will likely be assigned to $|h\rangle$ since the resonator response at the readout frequency is mostly flat for $|h\rangle$. The average assignment error for the four states is 7.(2)% [Fig. 4.12(b)] while the average qubit readout error is 1.(3)% [Fig. 4.12(c)]. Here, we assume that state preparation errors are small compared to assignment errors.

4.4.9 Leakage reduction for higher states

Although most common leakage mechanisms usually populate the second-excited state of the transmon, $|f\rangle$, some operations such as the measurement can leak into higher-excited states [18]. We observe the build-up of population in these higher states in the repeated parity-check experiment (Fig. 4). Figure 4.14 shows the fraction of total leakage to these higher states for the ancilla. Therefore, leakage reduction for higher states is necessary for ancillas. Similar to the leakage reduction mechanism that drives $|f\rangle \rightarrow |g\rangle$ [with effective coupling

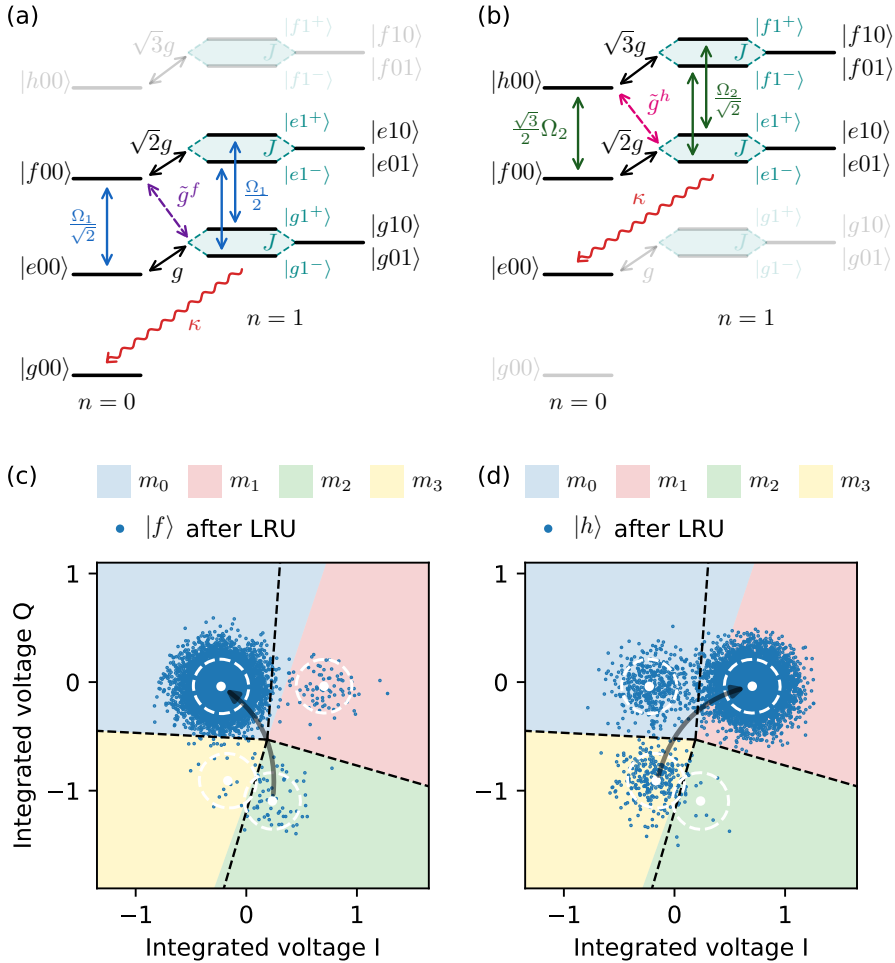


Figure 4.13: **Leakage reduction for $|f\rangle$ and $|h\rangle$.** (a, b) Transmon-resonator system level structure showing the relevant couplings for the f -LRU (a) and h -LRU (b). Each effective coupling, $\tilde{g}^f$ and $\tilde{g}^h$, is mediated by its respective drive Ω_1 and Ω_2 and transmon-resonator coupling g . (c, d) Readout data (2^{13} shots) of leakage states $|f\rangle$ (c) and $|h\rangle$ (d) after applying both LRU pulses simultaneously. The white dots and dashed circles show the mean and 3σ standard deviation obtained from fitting calibration data for each state.

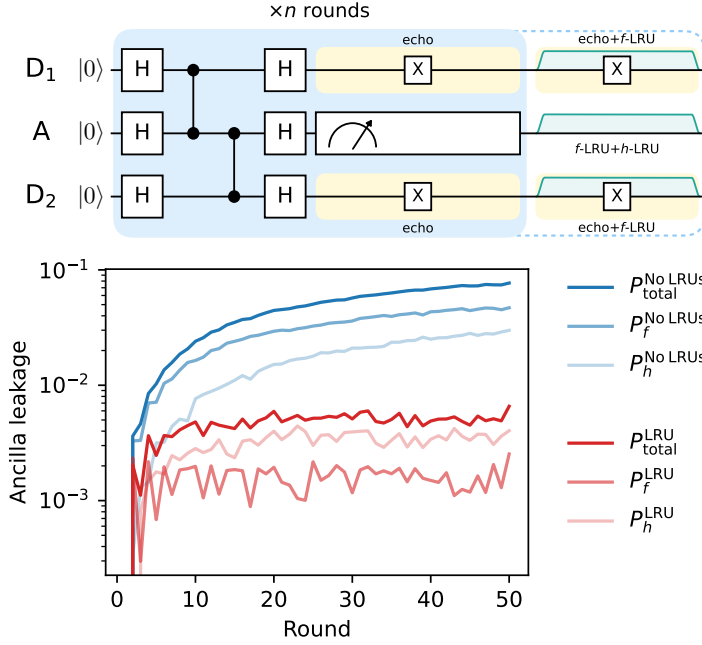


Figure 4.14: **Higher leakage states of the ancilla qubit.** Composition of ancilla leakage during the repeated stabilizer measurement of Figure 4. $P_{\text{total}} = P_f + P_h$ denotes the total leakage population without (blue) and with (red) LRUs.

$\tilde{g}^f$ in Fig. 4.13(a)], one can drive $|h\rangle \rightarrow |e\rangle$ (with effective coupling $\tilde{g}^h$ in Fig. 4.13(b)]. This transition can be induced much like the former, with an extra drive at frequency

$$\omega_{h00} - \omega_{e1\pm} \approx 2\omega_Q + 3\alpha - \omega_{R/P}, \quad (4.18)$$

2α below the f -LRU transition. The effective coupling for each LRU is given by (Eq. A35 of Ref. [72]):

$$\tilde{g}^f \approx \Omega \frac{g_{\text{eff}} \alpha}{\sqrt{2} \Delta (\Delta + \alpha)} \quad (4.19)$$

and

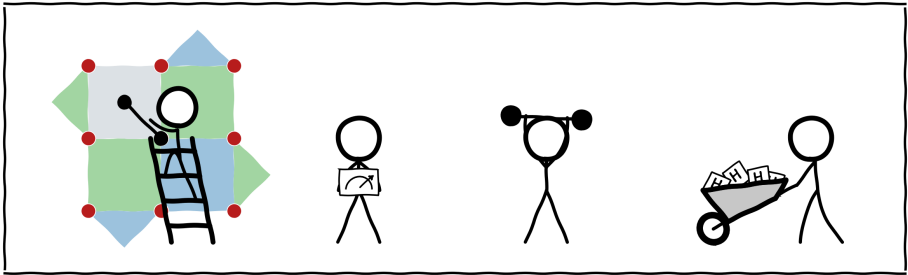
$$\tilde{g}^h \approx \Omega \frac{\sqrt{3} g_{\text{eff}} \alpha}{\sqrt{2} (\Delta + \alpha) (\Delta + 2\alpha)}, \quad (4.20)$$

where $\Delta = \omega_Q - \omega_{R/P}$ and $g_{\text{eff}} = g/\sqrt{2}$ (see section S3). As $\tilde{g}^h/\tilde{g}^f = \sqrt{3}\Delta/(\Delta + \alpha) > 1$, one should be able to drive $|h\rangle \rightarrow |e\rangle$ with comparable performance using similar drive amplitudes. We then have two LRU mechanisms, f -LRU and h -LRU, increasing seepage from $|f\rangle$ and $|h\rangle$, respectively. We drive both of these transitions simultaneously using two independent drives. Following the same calibration procedure shown in Fig. 2 for the f -LRU,

we tune up a pulse for the h -LRU. Figures 4.13(c) and 4.13(d) show readout data for states $|f\rangle$ and $|h\rangle$ after performing both LRUs simultaneously. The corresponding removal fraction for each state is $R^f = 99.2\%$ and $R^h = 96.1\%$ for $t_p = 220$ ns. Using this scheme, we can effectively reduce leakage in both states (Fig. 4.14). In particular, leakage in $|f\rangle$ is effectively kept under 0.2% , while that in $|h\rangle$ sits below 0.4% (red curves in Fig. 4.14). The former shows a flat curve and therefore corresponds to the L_1 of the cycle (similar to Fig. 4). The apparent remaining leakage in $|h\rangle$ could possibly be due to higher-excited states, which are naively assigned as $|h\rangle$ by the readout as they cannot be distinguished. One could potentially address these with additional drives. The general expression for the effective coupling of the transition targeting the m th leakage state (with $m_f = 0$, $m_h = 1$, etc.) is given by (Eq. A35 of Ref. [72])

$$\tilde{g}^m/\Omega \approx \frac{\sqrt{(m+1)(m+2)}g_{\text{eff}}\alpha}{2(\Delta + m\alpha)(\Delta + (m+1)\alpha)}, \quad (4.21)$$

which increases monotonically with m . Therefore, one could expect comparable performances for LRUs acting on higher states.



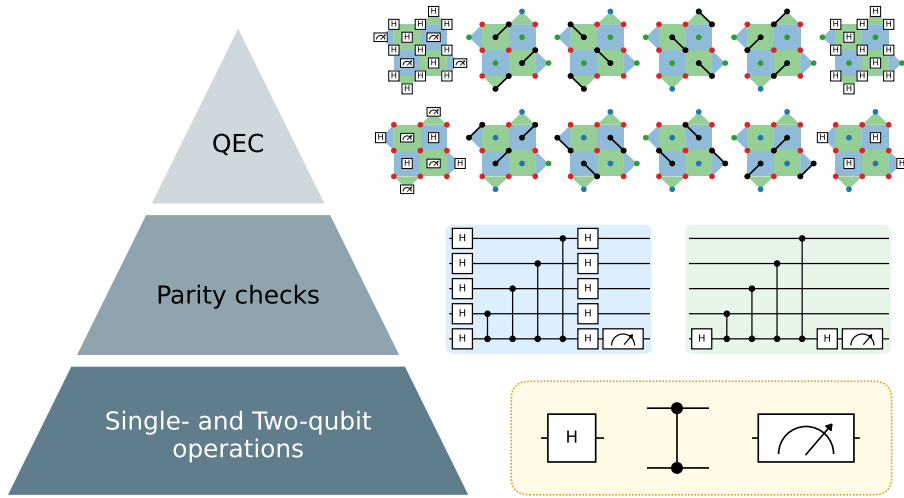


Figure 5.1: **The building blocks of quantum error correction.** We can think of building a quantum error correction code as pyramid of increasing complexity levels. At the bottom of this pyramid are the single- and two-qubit operations, followed by parity checks and the full quantum error correction cycle at the top. Because of crosstalk, it is useful to assess not only the individual gate operations that make up the lower level of the pyramid, but also benchmark these when put together in increasingly complex circuits as shown in the higher levels.

5.1 Introduction

Quantum error correction (QEC) provides an alternative to protect quantum information from errors due to decoherence and operational imperfections at the cost of redundancy. One way to do this, is by using stabilizer codes where the information is encoded across a multiple identical qubits. Such codes rely on the systematic procedure of measuring stabilizer observables, which are designed to discretize, detect, and correct errors without collapsing the encoded logical quantum state. This process forms the backbone of most QEC protocols, such as the surface code [62], enabling the identification of error syndromes that signal the occurrence of qubit errors. Parity checks (Fig. 5.1) measure the necessary observables to construct error syndromes. The precision with which parity checks are executed and interpreted (by means of decoding) dictates the overall performance of the quantum error correction scheme. As illustrated in Figure 5.1, to make a robust QEC implementation, it is crucial to thoroughly assess not only individual qubit operations but also the constituent parity checks that compose the stabilizer measurements. In current superconducting processors, crosstalk stands in the way of climbing this complexity pyramid. Therefore, it is important to develop

benchmarking procedures for all levels of this pyramid. An important aspect of understanding and improving QEC in hardware is the study of how physical error mechanisms manifest within stabilizer syndromes. Different error mechanisms — ranging from simple bit- or phase-flips to more complex leakage errors — leave distinct signatures in the stabilizer syndromes. These offer guidance in debugging and developing QEC experiments.

This chapter is structured to first address the benchmarking of parity check operations, assessing their performance and reliability in the context of QEC. We then proceed to analyze how physical qubit errors manifest within the syndromes, investigating how distinct mechanisms lead to correlated defect patterns and how they can be used to gauge logical performance. These collectively aim to provide a detailed examination of the critical components of QEC. Finally, we showcase quantum error correction experiments performed in a 17 qubit device.

5.2 Benchmarking quantum parity checks

Quantum parity checks measure the joint data-qubit observables that constitute the stabilizers of the QEC code. Therefore, the efficacy of QEC relies on the precise and reliable execution of parity checks. In chapter 2, we have discussed the calibration and benchmark of the individual qubit operations that make up the parity check circuit. While individual benchmarks are valuable indicators of overall performance, their collective behavior is often undermined by crosstalk, which can significantly alter performance. Given this complexity, it is very useful to have benchmarks for parity checks so that one can validate the collective performance of the underlying operations. To this end, we discuss three benchmarking procedures designed to evaluate different aspects of parity checks, each contributing to a multi-faceted view of collective operation performance:

5.2.1 Parity assignment fidelity

The first metric pertains to the ability of the ancilla qubit to discriminate parity between different states of the data qubits. One can assess this by preparing the data qubits in a collectively defined parity state, which is an eigenstate of the stabilizer to be measured. For Z -type stabilizers, these correspond to all computational states of the data qubits. Figure 5.2 illustrates this procedure for a weight-4 parity check. Here, we seek to evaluate how well the $ZZZZ$ parity is mapped onto the ancilla qubit state. To quantify this, we compute the parity assignment fidelity, averaged over the sixteen computational states of the data qubits (as shown in Fig. 5.2c). Additional insights can be gleaned by analyzing the relative assignment error between data-qubit states. Figure 5.2c indicates that states with a higher number of excited data qubits generally have a higher assignment error, which is expected due to increased susceptibility to relaxation. Higher errors common to states where a particular data qubit is excited could point to suboptimal two-qubit gate performance involving that qubit. Conversely, if higher errors are specific to states where two data qubits are excited, this could suggest

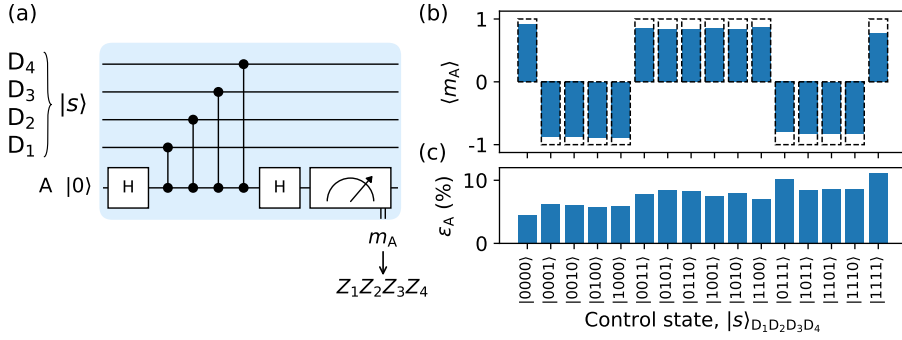


Figure 5.2: **Parity assignment fidelity measurement.** (a) Weight-4 Z -type parity check performed by the ancilla qubit A on the data qubit register $\{D_1, D_2, D_3, D_4\}$. The data qubit register is prepared in the computational state $|s\rangle$, whose Z -parity is mapped onto the ancilla qubit such that its ideal measurement outcome is $m_A = Z_1 Z_2 Z_3 Z_4$. (b) Average ancilla qubit outcome, $\langle m_A \rangle$, for all 2^4 computational states of the data qubits. The ideal and measured results are represented by the dashed wireframe and blue bars, respectively. (c) Parity assignment error, ϵ_A , for each computational state, ordered by the number of excited qubits. The average parity assignment fidelity is 92.3%.

a crosstalk mechanism. However, while this procedure is useful in evaluating the parity outcome, it does not account for errors that may reside in the data qubits post-measurement. These errors, which equally impact the efficacy of QEC, will be the focus of the next subsection.

5.2.2 Projection onto stabilizer subspace

The second procedure assesses how well the parity check projects the data qubits onto the stabilizer subspace. As shown in Figure 5.3a, this is accomplished by measuring an X -type stabilizer on a data qubit register initially in the state $|0\rangle^{\otimes 4}$. Using state tomography to reconstruct the data-qubit density matrix, we compare it to the ideal target states expected for each parity outcome, $\rho_{|\Phi+\rangle}$ and $\rho_{|\Phi-\rangle}$ (see Fig. 5.3b). The fidelity derived from this comparison serves as a direct indicator of the parity check's ability to project onto the intended stabilizer subspace, which will later define the codespace. Additionally, the reconstructed density matrices are particularly useful for detecting coherent errors in the data qubits resulting from calibration and crosstalk errors. For instance, single-qubit phase errors remaining after the parity check manifest as the complex phase of the off-diagonal elements of ρ , denoted as $\text{Arg}(\rho)$. Similarly, single-qubit phase errors occurring between single-qubit gates could also result in an X or Y rotation. Delaying the tomography measurements to after the ancilla measurement can also evaluate dephasing during this step.

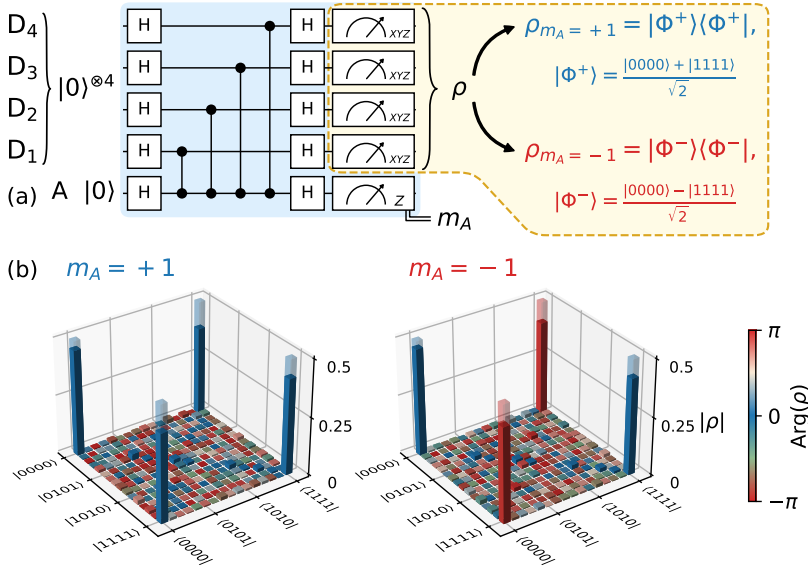


Figure 5.3: **Data qubit tomography measurement.** (a) A GHZ state is created in the data qubit register $\{D_1, D_2, D_3, D_4\}$ by measuring a weight-4 X -type parity check using the ancilla qubit A . The data qubits are initialized in the state $|0\rangle^{\otimes 4}$. Since this state is not an eigenstate of the parity check observable, $X^{\otimes 4}$, it will ideally be collapsed onto the GHZ state $|\Phi^+\rangle$ (even parity) or $|\Phi^-\rangle$ (odd parity), corresponding to the ancilla outcomes $m_A = +1$ and $m_A = -1$, respectively. The resulting state of the data qubits, ρ , is measured using state tomography conditioned on the ancilla qubit outcome. (b) Density matrices constructed for each ancilla qubit outcome, m_A . The fidelity to their ideal target states, $|\Phi^+\rangle$ and $|\Phi^-\rangle$, represented by the shaded wireframes, is 82.0% and 84.2%, respectively.

Up to this point, we have presented two complementary procedures: the first assesses the ancilla qubit's ability in assigning parity to the data qubits, while the second gauges the projection of data qubits within the stabilizer subspace. A comprehensive metric is now required to merge these two aspects, providing a holistic measure of the parity check. This will be the subject of the next subsection.

5.2.3 Repeatability

The final metric we evaluate is repeatability. This measures the probability of obtaining correlated parity outcomes from two consecutive parity checks (see Fig. 5.4a). Along with the parity assignment and data qubit projection errors evaluated in the previous methods, repeatability is also sensitive to additional error mechanisms that these methods do not detect. One such mechanism is the dephasing of data qubits during the ancilla qubit measurement, which could

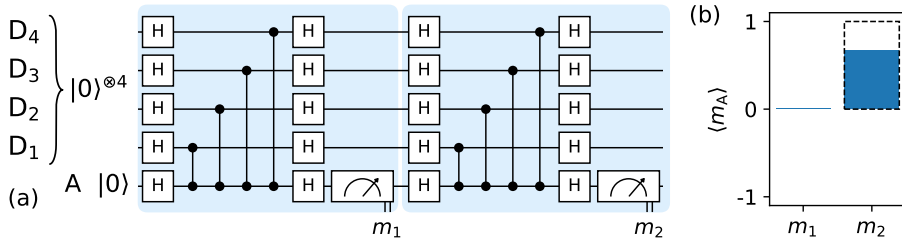


Figure 5.4: **Parity check repeatability measurement.** (a) Two consecutive X -type parity checks are performed on the data qubits $\{D_1, D_2, D_3, D_4\}$ using the ancilla qubit A . The initial state of the data qubits, $|0\rangle^{\otimes 4}$, is projected onto an even or odd GHZ state after first parity check. Ideally, this occurs with equal probability based measurement outcome m_1 . In the subsequent parity check, the data qubits state, now an eigenstate of the stabilizer $XXXX$, should ideally be preserved and the second measured outcome, m_2 , should reflect the same parity. Since the behavior of the parity check is to maintain (flip) the ancilla state if parity is even (odd), then ideally $m_2 = +1$. (b) Experimental average measurement outcomes of m_1 and m_2 . The calculated repeatability for this measurement is 80.4%.

5

arise from readout crosstalk. Ancilla qubit errors during readout due to the non-quantum non-demolition (QND) nature of the measurement will also affect repeatability. Repeatability is an excellent indicator of the overall performance of the parity check as it reflects the impact of all types of error: data qubit, ancilla qubit, and assignment errors. Therefore, it serves as a gauge for the incidence of general errors in the parity check and a valuable tool to validate calibration at a collective level in the device.

Through these three benchmarking procedures, we can obtain a comprehensive understanding of the performance of parity checks. These metrics not only reflect the individual operations' performance but also reveal the nuanced dynamics of their collective behavior, which is important when performing any quantum algorithm.

5.3 Signature of errors in repeated stabilizer measurements

After benchmarking the building blocks of quantum error correction — parity checks — we now turn our attention to understanding the signature of physical qubit errors as they occur in stabilizer measurements. Understanding how physical qubit errors manifest in the outcomes of stabilizer measurements is critical for quantum error correction. A logical quantum state is corrected within a specific subspace, defined by a set of stabilizer observables and their corresponding eigenvalues — or parities. Errors in the constituent physical qubits of the QEC code lead to changes in the parity of stabilizers, known as defects. Therefore, to correct errors in the logical qubit, we must understand the signature of physical qubit errors in the syndrome of defects. Figure 5.5 shows how to compute the syndrome of defects, σ_A , from measured par-

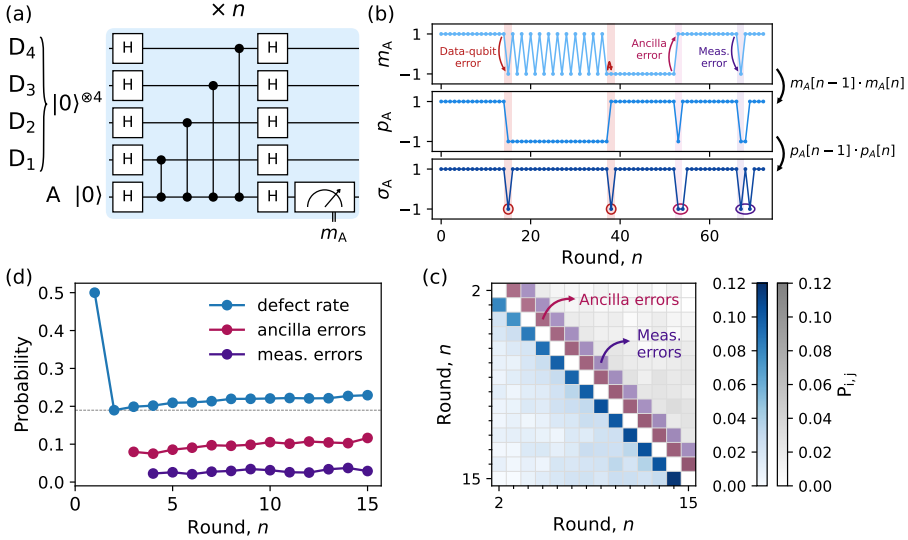


Figure 5.5: **Physical qubit error syndromes in repeated parity checks.** (a) Repeated X -type parity checks are performed on the data qubits $\{D_1, D_2, D_3, D_4\}$ using the ancilla qubit A . The measured outcomes, m_A , can be used to signal errors. (b) An example set of measurement outcomes, m_A , is converted into parity outcomes, p_A , and syndromes, σ_A . Physical qubit errors manifest into different behaviors of these quantities and are best understood through their σ_A signature. Data-qubit errors, such as a phase flip occurring between rounds, alter the X -parity of the data-qubit state, causing m_A to switch from constant (indicating even parity) to alternating (indicating odd parity) values. Consequently, these errors result in a single defect ($\sigma_A = -1$) detected in A . Ancilla-qubit errors, such as a bit flip between rounds, flip m_A and cause two consecutive defects. Readout errors in the ancilla cause m_A to invert twice, thereby producing two defects separated by one rounds. (c) P_{ij} matrix for a weight-4 parity check showing the correlation between defects occurring across different rounds. The relative frequency of ancilla qubit flips and measurement errors is manifested in the matrix's first and second off-diagonals, respectively, due to their distinct σ_A signatures. (d) Probability of a defect pattern occurring in each round of the parity check. The defect rate (blue) shows the probability of a single defect occurring in any given round, i.e., $P(\sigma_A = -1)$; ancilla errors (pink) indicate the probability of two consecutive defects occurring; measurement errors (purple) show the probability of two consecutive defects occurring, separated by one round.

ities, p_A , derived from the ancilla measurement outcomes, m_A , of a parity check. Each type of physical qubit error leaves a characteristic pattern of defects in the syndrome, revealing the nature of the underlying error mechanism. A decoder will try to decipher this information and deduce its implications for the logical qubit state, aiming to correct errors without disrupting the encoded information. From an experimental standpoint, analyzing the frequency of these error signatures provides insight into the prevailing error mechanisms. This information is valuable, not only for validating the efficacy of current calibration techniques *in situ* but also for identifying potential crosstalk issues that may occur. Furthermore, recognizing the most frequent sources of error enables targeted improvements in both calibration protocols and hardware design, guiding the development of better QEC systems.

In this section, we will study the signature of different physical error sources first in the context of individual stabilizer measurements and then in multiple stabilizer measurements. Although the examples shown here will use X -type stabilizers, the conclusions are analogous for Z -type stabilizers as well.

5

5.3.1 Ancilla qubit and measurement errors

We start by considering repeated measurements of a weight-4 $XXXX$ stabilizer, as illustrated in Figure 5.5a. After an initial parity check projects the data qubits into a specific parity subspace, any subsequent errors on the data qubits that do not commute with the stabilizer observable will change the state's parity. Errors such as Z or Y will trigger a defect ($\sigma_A = -1$) in the syndrome, as shown in Figure 5.5b. Errors in the ancilla qubit, such as relaxation during readout or a phase-flip between CZ gates, will lead to a perceived momentary change in parity, resulting in two consecutive defects in the syndrome. If an error occurs instead in the assignment of the ancilla outcome, this will again result in two defects, now separated by one round. Note that both the type of error and the particular point in the circuit where it occurs determine the defect pattern. A useful metric that quantifies the incidence of general errors in a parity check is the average defect rate. Figure 5.5d shows this rate measured for a weight-4 stabilizer across multiple rounds of parity checks. This rate, which is closely related to the repeatability metric discussed in the previous section, indicates the probability of a defect occurring and its variation over successive rounds. While the defect rate aggregates the overall error rate in a parity check, it is possible to differentiate between assignment and ancilla qubit errors from the overall error rate. This distinction is based on the fact that assignment and ancilla qubit errors lead to the occurrence of two defects in the syndrome of the stabilizer. By examining the correlations between defects across different rounds, we can quantify the probability of these errors separately. Assuming that all correlations are pairwise and that errors are sparse enough in time, the probability of simultaneously triggering two defects at rounds i and j , denoted P_{ij} , is given by [46]:

$$P_{ij} = \frac{\langle \sigma_i \sigma_j \rangle - \langle \sigma_i \rangle \langle \sigma_j \rangle}{4 \langle \sigma_i \rangle \langle \sigma_j \rangle}. \quad (5.1)$$

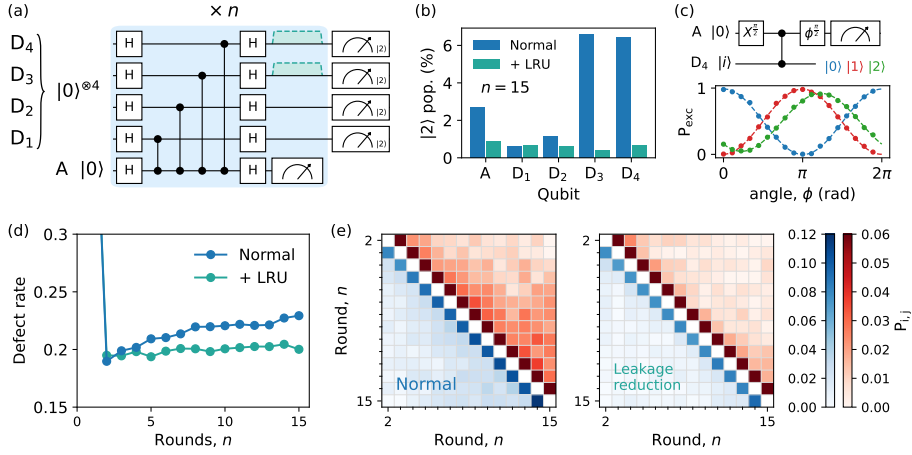


Figure 5.6: **Impact of leakage in repeated parity checks.** (a) Repeated X -type parity checks are performed on the data qubits $\{D_1, D_2, D_3, D_4\}$ using the ancilla qubit A. (b) Multi-level readout is used to measure leakage in the data qubits after $n = 15$ rounds of parity checks. Leakage reduction units (cyan pulse) are employed on the data qubits most prone to leakage. (c) Ramsey experiment used to measure the phase acquired by the ancilla qubit, A, during a CZ gate with data-qubit D_4 for computational states $|0\rangle$ and $|1\rangle$ and leakage state $|2\rangle$. This phase is 0 and 180 and 139 degrees for $|0\rangle$, $|1\rangle$ and $|2\rangle$, respectively. (d) Defect rate at round n for the parity check circuit (blue) with the addition of LRUs (cyan). (e) P_{ij} matrix computed from the ancilla qubit syndrome, m_A , for the parity check circuit (Normal) with leakage reduction (LRU). Full and saturated scales are shown in blue and red, respectively. Long-term correlation between defects is substantially reduced when using leakage reduction.

Figure 5.5c displays the P_{ij} matrix for up to 15 rounds. A couple of points to note: the P_{ij} matrix is symmetric, so its typically represented using two color scales to emphasize certain features [46]; additionally, correlations including the first round are omitted since the initial syndrome measurement is random. Errors in the ancilla qubit create a correlation between a defect in round i and another in round $i + 1$, which are then visible on the first off-diagonal of the matrix. Similarly, measurement errors, which correlate defects separated by one round, manifest along the second off-diagonal. We can see how these specific error-induced defect patterns compare to the overall defect rate in Figure 5.5d. Such analysis is a valuable tool for gaining a clearer understanding of the sources of errors in the parity check.

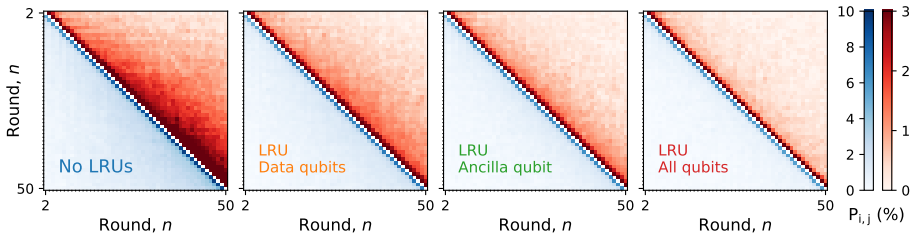


Figure 5.7: **Correlation of defects in repeated stabilizer measurements using leakage reduction units [106].** P_{ij} for the four experiments reported in Figure 4.4 of Chapter 4.

5.3.2 Leakage errors

As observed in our experiments, the defect rate tends to increase with the number of rounds, suggesting a growing incidence of errors. This trend is threatening for quantum error correction (QEC), as fault-tolerance requires error rates that are not only low but also stable. Such a pattern of increasing error rates over time has been consistently reported in the literature [38, 71, 99, 106, 107], including in Chapter 4 of this thesis. This phenomenon is commonly attributed to the accumulation of leakage in the transmon. Leakage, being a non-Pauli error, poses a unique challenge in that it does not produce a consistent pattern of defects in the stabilizer syndrome. To understand the impact of leakage on the stabilizer syndrome, we perform repeated measurements of a weight-4 stabilizer, applying leakage reduction units [106] to the qubits most affected by leakage, as shown in Figure 5.6a. Figure 5.6b shows there is reduction in leakage population across the board, even though leakage reduction was only directly applied to qubits D_3 and D_4 . Predicting the precise effect of leakage on the syndrome is challenging, particularly when it involves the ancilla qubit [70]. Since we use binary measurement outcomes, the influence of leakage on the syndrome depends on the technical details of the readout scheme. Specifically, how leakage is classified in the heterodyne measurement apparatus [13]. For data-qubit leakage, its effect on the ancilla must be taken into account. A data qubit in a leaked state will impart a conditional phase on the ancilla qubit during the CZ interaction [99]. This phase, in turn, determines the likelihood of a defect occurring. Figure 5.6c shows the conditional phase imparted by D_3 , measured as an example. Any deviation from 0 or 180 degrees will result in a non-zero probability of defects occurring. While the exact pattern of leakage in the syndrome can be complex, one consistent effect is the presence of long-term defect correlations. This is attributed to the fact that a leaked qubit may remain in its leaked state for several rounds before it returns back to the computational state. Such behavior leads to persistent correlations that extend over multiple rounds, becoming apparent in the off-diagonal elements of the P_{ij} matrix. These long-term correlations are significantly reduced when using leakage reduction as shown in Figure 5.6e. This effect is even more pronounced in the experiment reported in Chapter 4

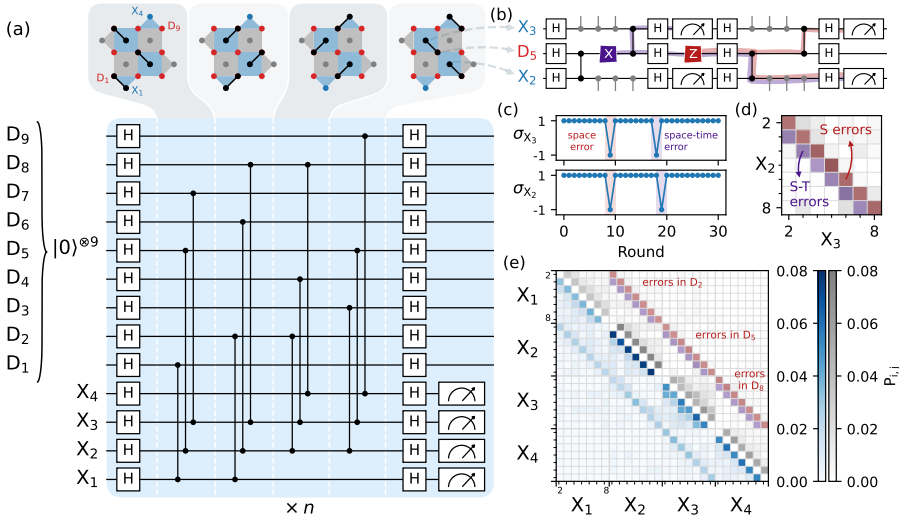


Figure 5.8: Data-qubit errors in repeated parity checks. (a) Repeated X -type parity checks in a distance-3 surface code, which include two weight-2 and two weight-4 stabilizers measured by ancilla qubits X_1 , X_2 , X_3 , and X_4 . (b) The propagation of errors through the circuit, involving data qubits at the intersection of stabilizers, gives rise to two types of error syndromes: space and space-time errors. A phase flip occurring in D_5 between rounds (indicated in red) results in simultaneous defects in the adjacent ancilla qubits X_2 and X_3 [space error, depicted in (c)]. A bit flip occurring between consecutive two-qubit gates (indicated in purple) leads to defects in X_3 and X_2 across consecutive rounds [space-time error, depicted in (c)]. (d) The P_{ij} matrix, correlating defects in X_2 and X_3 over different rounds, reveals space errors along the main diagonal and space-time errors along the first off-diagonal. (e) The P_{ij} matrix for all ancilla qubits over 10 rounds of parity checks shows the magnitude of space and space-time errors detected in data qubits D_2 , D_5 , and D_8 on the shaded diagonals.

Figure 4.4. Here, LRUs were progressively added on all qubits and performed for many more rounds. Figure 5.7 shows P_{ij} for each variant of the experiment, where we find the gradual reduction of long-term correlations. Reducing leakage at each round prevents its build-up in the transmon, which, in turn, contributes to stabilizing the occurrence of errors over time. Hence, leakage reduction also leads to a more uniform defect rate, as shown in Figure 5.6d. Complementing the observations made in Chapter 4, this emphasizes the critical role leakage reduction for effective quantum error correction.

5.3.3 Data qubit errors

In the previous sections, we established that errors in a data qubit trigger a single defect in the stabilizer syndrome. While that alone flags an error occurrence, this information is not sufficient to identify the data qubit in which the fault occurred. Spatial resolution of individual data qubit errors becomes possible when a qubit is shared between two independent stabilizers. In such cases, an error in a data qubit will produce a pair of defects, one in each stabilizer. By analyzing the correlations between defects in the two stabilizers, we can precisely quantify the error probability for the individual qubit in question. An example of this is shown in Figure 5.8, where repeated X -type stabilizers are measured in a distance-3 surface code. This code configuration, which includes four X -type stabilizers intersecting different three data qubits, is depicted in Fig. 5.8a. In this experiment, the initial state of the data qubits, $|0\rangle^{\otimes 9}$, is projected into the surface code state, $|0_L\rangle$, according to Born's rule,

$$|0_L\rangle = \frac{1}{4} \prod_i (I + m_i S_i^X) |0\rangle^{\otimes 9}, \quad (5.2)$$

where the product runs over all X -type stabilizers of the code, S_i^X , and its respective measurement outcome, m_i . Data qubit phase errors occurring in this state will trigger defects in the stabilizer syndrome. One of particular interest cases is the central data qubit in the bulk of the code, D_5 . As illustrated in Figure 5.8b, a Z -error between parity check rounds will trigger two simultaneous defects in the stabilizer syndromes of X_2 and X_3 (see Fig. 5.8c). These are referred to as space-like errors. Conversely, if an X -error occurs between the two-qubit gates of the parity checks (indicated in purple in Fig. 5.8b), it triggers two consecutive defects, one in each stabilizer, as depicted in Fig. 5.8c. These are called space-time-like errors. The probability of occurrence of these errors can be measured by calculating P_{ij} between rounds of the two stabilizers. In this matrix (Fig. 5.8d), space-like errors are represented on the diagonal, while space-time-like errors appear on the first off-diagonal. We find that space errors are more common than space-time errors. This observation is consistent with the time windows for each error type: the window for space-time errors is much shorter (approximately 120 ns) compared to that for space-like errors (approximately 500 ns). By calculating P_{ij} between all stabilizers, as shown in Figure 5.8e, one reveals data qubit errors in D_2 and D_8 as well.

These techniques provide a holistic method for assessing the performance of a quantum error correction code without the need to directly compute logical error rates, which would typically require a decoder. This approach is particularly beneficial in experimental contexts. It enables the identification of predominant error types and their sources, while also exposing crosstalk. Such insights are important to develop calibration methods and making necessary adjustments in hardware design to effectively mitigate these errors. Furthermore, this information is also valuable in the context of decoders. For example, in a minimum-weight-perfect-matching (MWPM) decoder, the information about defect correlations can be utilized to calibrate the weights assigned to edges in the decoding graph. This calibration is crucial for

optimizing the performance of error correction. It is important to emphasize that the validity of these metrics relies on the assumption that errors are sufficiently low, ensuring that defect correlations are representative of the actual frequency of error occurrences. Only under this sparsity condition can we reliably infer error rates from the observed defect patterns.

5.4 Correcting bit flips in a distance 3 surface code

In this section, we investigate bit flip errors within the framework of a distance-3 surface code. This particular experiment is designed to measure only the Z -type stabilizers of the code. This approach allows us to circumvent some of the problematic gate performance issues by utilizing 13 out of the 17 qubits in our device. Our initial objective is to study the incidence and nature of physical qubit errors that occur within this specific scheme. By doing so, we aim to understand the predominant error mechanisms in the device. Following this analysis, we apply this scheme to correct bit-flip errors in computational states of the data qubits. By averaging over different computational states, we estimate a logical error rate for this code to evaluate the efficacy of error correction on the encoded logical information, similarly to a logical state in the surface code.

5.4.1 Assessing physical qubit errors *in situ*

As studied in the previous sections, we can gauge the general incidence of physical qubit errors in a logical qubit by looking for defects (i.e., changes in parity) in the syndrome of stabilizers of the code. In this simplified version of the surface code, we measure four Z -type stabilizers using two weight-2 and two weight-4 parity checks. To obtain a representative assessment of the error incidence in this code, we conduct repeated stabilizer measurements on the state $|+_{\text{L}}\rangle = |0_{\text{L}}\rangle + |1_{\text{L}}\rangle$. This state is the simplest non-trivial state within the logical subspace of this surface code that can be prepared using only the Z -type stabilizers. The state is prepared by initializing the data qubit register in $|+\rangle^{\otimes 9}$. Since this product state is not an eigenstate of the stabilizers, an initial round of stabilizer measurements will randomly project the data qubits onto the state:

$$|+_{\text{L}}\rangle = \frac{1}{4} \prod_i (I + m_i S_i^Z) |+\rangle^{\otimes 9}, \quad (5.3)$$

where the product is taken over all Z -type stabilizers of the code, S_i^Z , and its respective measurement outcome, m_i . The defect rate for each stabilizer, $\{Z_1, Z_2, Z_3, Z_4\}$, is shown in Figure 5.9a. We observe that these rates are divided into two distinct groups: approximately 21% and 14% for weight-4 and weight-2 stabilizers, respectively. This discrepancy is expected due to the higher number of error channels in a weight-4 parity check, which involves five qubits, as compared to a weight-2 check that involves only three qubits. We can differentiate between the various error mechanisms contributing to the overall defect rate

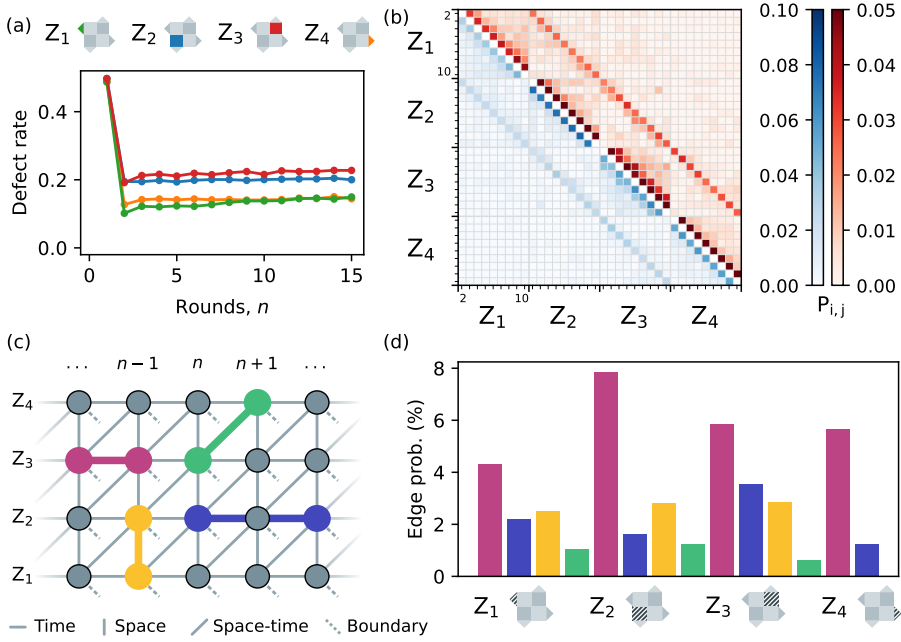


Figure 5.9: Space-time graph of defects. Repeated Z -type stabilizer measurements in a distance-3 surface code. The data qubits are initially prepared in the state $|+\rangle^{\otimes 9}$ and subsequently projected into the logical state $|+_L\rangle$. (a) Defect rate for each of the four Z -type stabilizers is presented. (b) P_{ij} matrix showing the correlation between defects in the stabilizer syndrome. (c) Space-time stabilizer defect graph is depicted, with nodes representing the syndrome outcome of a stabilizer in space (vertically) and time (horizontally). The graph's edges denote physical qubit errors: time-like edges of length 1 and 2 indicate ancilla qubit and measurement errors, respectively (see Fig.5.4); space-like and space-time-like edges denote errors in a data qubit that interacts with two stabilizers (as shown in Fig.5.8); boundary edges indicate errors in a data qubit measured by a single stabilizer. (d) Measured probability of different edges in the graph, averaged across all rounds of the experiment.

by examining the correlation between defects across different rounds and stabilizers. To accomplish this, we calculate P_{ij} [46, 108] (Eq. 5.1) between all four stabilizers (Fig. 5.9b). To associate specific error mechanisms with observed correlations, it is convenient to look at defects in a space-time graph as depicted in Figure 5.9c. Each node in this graph represents a syndrome measurement of a stabilizer in space and in time. Physical qubit errors are expressed as edges in this graph, connecting either two defects or a single defect to the boundary. As described in Section 5.3.1, ancilla qubit and measurement errors manifest as time-like edges in the graph separated by one and two rounds, respectively. Data qubit errors are represented as space or space-time-like edges between defects in different stabilizers (see Section 5.3.3). Errors occurring in data qubits at the boundary of the code, which are connected to a single stabilizer, lead to a single defect and are denoted as boundary edges. The P_{ij} matrix contains the probability of each type of edge in the space-time graph, except for boundary edges. Figure 5.9d shows these probabilities, averaged over time, for the aforementioned edges. From this analysis, we find that ancilla qubit errors (pink) are the predominant error in our code. This seems to be particularly problematic for ancilla qubit Z_2 , which we link to a TLS issue (see Sec. 2.4). Data qubit errors (yellow and green) are the second most frequent source of errors, with measurement errors (purple) being the least prevalent. The high rate of measurement errors in ancilla qubit Z_3 suggests low readout assignment fidelity.

We note that additional error mechanisms can lead to other types of edges in the space-time graph. For instance, in a conventional surface code that includes both Z - and X -type stabilizers, hook errors [65] manifest as space-like edges of length two. Furthermore, more complex hyperedges [90, 109] can arise from leakage and crosstalk mechanisms, potentially resulting in higher degree correlations that encompass more than two defects. Leakage, which we have already observed leading to long-term time-like correlations (see Section 5.3.2), is an example of this as long lived leakage can potentially lead to many defects. It would be great to add the probability of boundary edges occurring into this analysis. Doing so would provide further insight into the remaining data qubit errors thus painting a more complete picture of the statistics of physical errors. To achieve this, one would need to accurately calculate the probability of all other error mechanisms triggering a defect at a specific node [46, 108]. However, this task is considerably complicated by the presence of hyperedges, as these higher degree correlations are challenging to predict in practice [109].

This *in situ* analysis of errors is useful not just for benchmarking and calibration, but are also the foundation for most decoders used in QEC. As we have seen in this section, the problem of decoding syndromes can be viewed as a problem of matching defects on a graph. This problem is commonly solved using a minimum-weight-perfect-matching algorithm (MWPM). MWPM uses these probabilities to infer the most likely set of errors for a given syndrome. Therefore, accurately identifying and quantifying physical error probabilities is key for improving the efficiency of the decoder and, in turn, reducing the logical error rate.

5.4.2 Logical error rate

Until this point, we have focused on the study of stabilizer measurements and how they can be used to gauge physical qubit errors. We will now use repeated stabilizer measurements to correct errors in the surface code. The correction of specific physical qubit errors (bit-flips or phase-flips) has been experimentally demonstrated in repetition codes, as reported in several studies [45–47, 49–51]. More recently, there have been significant advancements in the repeated quantum error correction of fully protected logical qubits [38, 88–90]. Particularly noteworthy is the recent breakthrough in reducing the logical error rate of a surface code by increasing the distance of the code [39]. In such experiments, the logical error rate is typically the main figure of merit. This metric quantifies the errors incurred by the logical qubit per round of error correction. To determine this rate, one typically initializes data qubits of the code in a logical state and measures the fidelity of the encoded information over a series of successive stabilizer measurements. We will emulate this by using repeated Z -type stabilizer measurements in the surface code to correct for bit-flip errors in the computational states of data qubits. Since computational basis states are eigenstates of both the Z -type stabilizers S^Z and the logical operator Z_L of the surface code, they can be used to store logical information. However, this information will only be protected from X errors, i.e., bit-flips, similarly to a repetition code.

Similarly to a logical state in the surface code, we will use computational states that comprise the superposition of a logical state. For this purpose we choose the 16 states of the superposition $|0_L\rangle$ such that,

$$|0_L\rangle = \frac{1}{4} \prod_i (I + S_i^X) |0\rangle^{\otimes 9}, \quad (5.4)$$

where the product runs over all the X -type stabilizers of the code S_i^X . The set of states described is represented schematically in Figure 5.10a. After preparing one of these states, $|s\rangle$, we will then perform n successive stabilizer measurements from which we construct the syndrome of stabilizers used to decode errors. To average the error rate over the logical subspace, we flip the logical state at the end of each round using an X_L transversal gate (Blue dashed box in Fig. 5.10a). A final measurement of all data qubits is used to measure the logical operator Z_L outcome and construct a final syndrome measurement (Fig. 5.10a).

To decode the syndrome measurements, we use a MWPM decoder to infer the Pauli frame updates required to correct errors. After applying the corrections, we compute the average $\langle Z_L \rangle$ for the different n rounds of the experiment and for each input state $|s\rangle$. From this average, we calculate the logical fidelity, F_L , using,

$$F_L = \frac{1 + |\langle Z_L \rangle|}{2}. \quad (5.5)$$

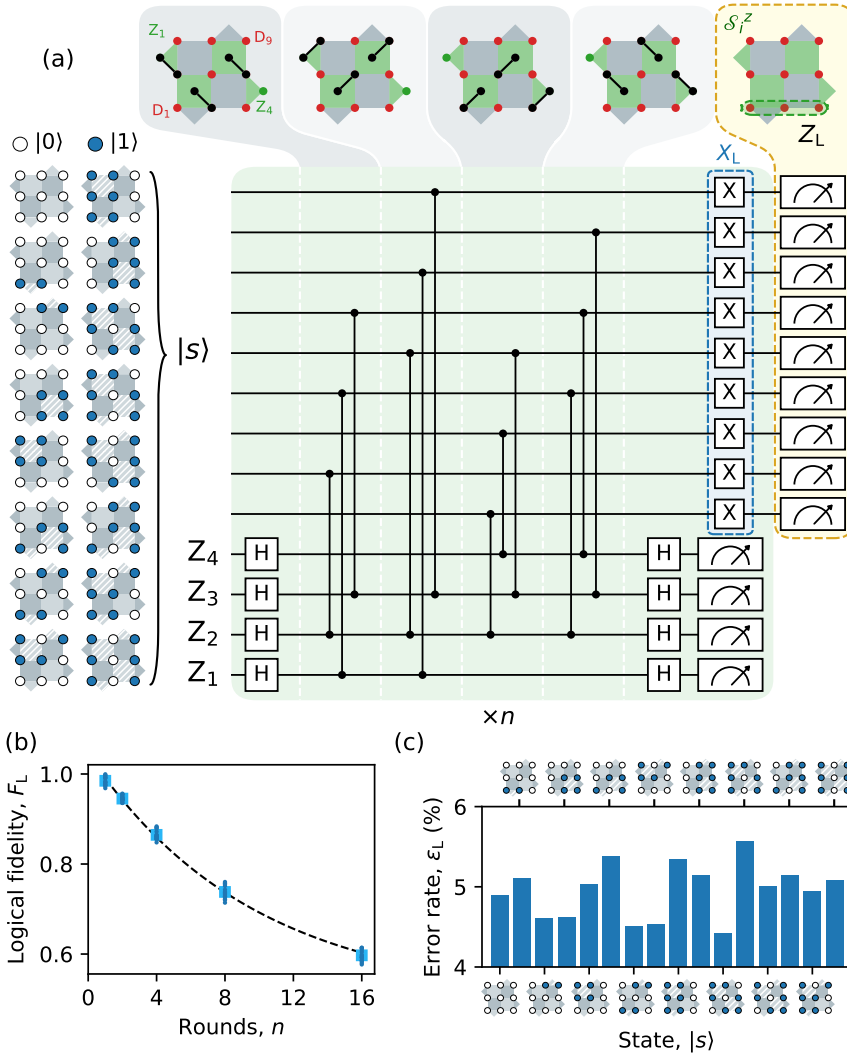


Figure 5.10: **Correction of Bit-Flips in the Distance-3 Surface Code.** (a) Error correction circuit. Data qubits are initially prepared in one of the computational states, $|s\rangle$, which are part of the superposition comprising the state $|0_L\rangle$ (see Eq. 5.4). This is followed by a sequence of n repeated Z -type stabilizer measurements. A final measurement of the data qubits is performed to determine $\langle Z_L \rangle$. (b) Logical error calculated using the average $\langle Z_L \rangle$ versus the number of parity check rounds, n , for each of the 16 states (dots) and for the average (square). The dashed curve shows fit to a model used to estimate the logical error rate, which is found to be 3.9% per round. (c) Logical error rate extracted for each individual computational state.

The results are shown in Figure 5.10b for the individual states, $|s\rangle$, as well as for the average of them. To find the logical error rate, ϵ_L , we fit the resulting data to the model [83],

$$F_L(n) = \frac{1}{2} \left[1 + (1 - 2\epsilon_L)^{(n-n_0)} \right], \quad (5.6)$$

where ϵ_L denotes the logical error rate, and n_0 accounts for the delay in the onset of exponential decay observed for small n [83]. We do this for F_L averaged accross all states and find a logical error rate, ϵ_L , of 4.(9)% (Fig. 5.10b). Alternatively, we can fit F_L for each individual input state. The resulting error rates (Fig. 5.10c) show a standard deviation of 0.3% which point to a uniform error rate for accross all states, $|s\rangle$.

6.1 Outlook

In this chapter, I will discuss the results of this thesis and contextualize them within the current literature. I will also give my perspective on the future directions of research. The field of superconducting qubits has seen remarkable progress during the course of my Ph.D., propelled by both small-scale efforts in academia and large-scale industry initiatives.

The first experimental work on quantum error correction dates back to the late 90s in nuclear magnetic resonance (NMR) systems [110], and it has since branched out to a variety of physical platforms including trapped ions [53, 54, 88, 91, 111–114], photonics [115–118], nitrogen-vacancy centers in diamond [50, 119–121], neutral atom arrays [122, 123], and superconducting circuits [38, 39, 45, 49, 57, 58, 63, 89, 90, 124–126]. In this discussion, my focus will be on superconducting qubits. While gate errors in this platform were nearing the threshold required for quantum error correction [8], the first demonstrations of ancilla-based parity checks, required for stabilizer codes, were being realized [66, 67]. Simultaneously, the first quantum error detection experiments were conducted for bit-flip [49, 124, 127] and arbitrary errors [125] in a single round, later extending to multiple rounds in repetition codes [45]. Advancements in dispersive readout [41, 77, 79, 128–130] enabled deeper multi-round experiments such as the use of repeated stabilizer measurements for correcting arbitrary quantum errors in entangled Bell states [57, 58]. Shortly thereafter, the field began to witness experimental implementations in the surface code [62, 131–133]. This code gained notoriety due to its moderately high error threshold ($\approx 1\%$) achieved using nearest-neighbor qubit connectivity in a 2D lattice making it suitable for monolithic architectures like circuit QED [1]. The first of these was repeated error detection in the smallest, distance $d = 2$, surface code [46, 63, 126]. The results of Chapter 3 are among these. Here, similar to previous work [46, 63] we performed repeated error detection. Additionally, we studied various techniques to initialize, measure and perform gates in this logical qubit. In particular, we studied fault-tolerant and non-fault-tolerant implementations and their fidelity in the logical subspace after error detection. Our gate set was comprised of transversal Clifford gates as well as non-Clifford gates like the T gate which allow universal control of the single logical qubit. Through simulation, we learned that conventional error mechanisms were not sufficient to explain the observed results. Together with evidence of accumulating leakage in the readout data, this suggested that leakage was playing a substantial role in the errors detected (and thus discarded) in the experiment. This was concerning, as such errors would have to be dealt with in future error correction experiments. This put leakage on our radar which later led to the experiments in Chapter 4.

After this experiment, we started working on the larger Surface-17 devices with the goal of quantum error correction. The larger scale of these devices meant more automation was required to handle calibration. In particular two-qubit gates which grew from 8 in the old 7 qubit device to 24 in the new 17 qubit device. This prompted the developments reported in Chapter 2. Here, the goal was to develop fast and reliable hands-off calibration and benchmark

routines that could be run on a daily basis. An outstanding challenge in this domain is having to contend with interacting spurious two-level systems [30–33, 134]. This issue has been reported in the literature as one of the main limiting factors to performance [38, 39, 134]. As shown in this chapter, most calibration procedures are straightforward until a TLS enters the picture. Their unstable and unpredictable nature adds to the complexity of this problem and posed the most severe challenge in the experiments reported in Chapter 5.

As the first repeated quantum error correction experiments started to emerge [38, 39, 88–90], the subject of leakage in superconducting qubits gained prominence [19, 70–72, 99, 106, 107] as one of the main threats to fault tolerance in these systems. In particular, leakage reduction units (LRUs) [85, 86, 93, 94] became a fundamental part of quantum error correction experiments [39, 46, 71, 99]. Among these experimental realizations was chapter Chapter 4. This experiment followed from the theory proposal of Battistel *et al* [72], which included an all-microwave leakage reduction technique. This technique set itself apart from previous implementations [71, 99] because it has benign impact on the qubit subspace making it suitable for simultaneous use in both data and ancilla qubits in an error correction cycle. Additionally, it can be adapted to reduce leakage in higher states as well, which proved to be important in our device due to measurement-induced transitions to higher states of the transmon [18, 19]. We implemented this LRU accross multiple qubits in our device spanning different frequency configurations demonstrating both its parameter flexibility and ease of calibration which are important for implementation accross large scale devices. Then, we showcased its use in a repeated stabilizer measurement where the LRU reduced leakage by an order of magnitude accross all qubits. Similar to previous work [71], we verify that reducing leakage reduces the overall detection of physical qubit errors as well as long-term correlations between them (Fig. 5.7). Here, we learned that a significant part of the leakage observed was found in state $|3\rangle$ and possibly above. Such superleakage could arise from transport of leakage in two-qubit gates (where levels in higher excitation manifolds can become resonant [70, 71]) or induced by the readout [18, 19]. To handle this, we added an additional drive to simultaneously reduce leakage in $|2\rangle$ and $|3\rangle$.

As we progressed with calibrating Surface-17 devices for quantum error correction, the increasing complexity of the circuits exposed crosstalk errors that did not manifest or were benevolent in the experiments of Chapter 3. Such errors are sometimes challenging to identify due to their convolution with other operations in the experiment. This prompted the need to go beyond the benchmarking of individual operations (shown in Chapter 2) to the techniques discussed in Chapter 5. Here, the aim is to establish benchmarking protocols that are practical in an experimental context and help to identify and quantify error mechanisms *in situ*. This helps gradually increase the scale of our experiments from individual operations, to parity checks, and ultimately to full quantum error correction. In this context, we start by establishing metrics to benchmark stabilizer measurements. Following this, we progress to look at how different physical errors manifest in repeated stabilizer measurements. Finally, we look

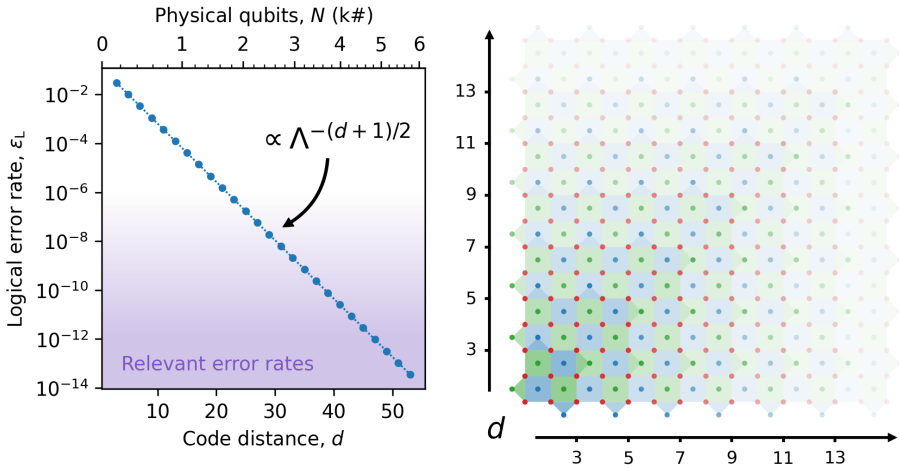


Figure 6.1: **Scaling of the surface code for useful applications.** (a) Logical error rate of a surface code logical qubit versus the code distance, d , and corresponding number physical qubits, N . The exponential suppression factor used for this estimate is $\Lambda = 3$. The shaded area denotes the error regimes relevant for useful large scale applications [62, 140–143]. (b) Illustration of surface code layout for increasing code distance, d .

6

at how we can use the aforementioned techniques to correct for bit-flip errors in a surface code. Our ultimate goal of full QEC in this experiment remains to be achieved. For a long time, our progress was hindered by TLSs resonant with qubits at the sweetspot. However, as discussed in Section 2.4.3, we have since addressed this issue by using unipolar gates. Future challenges in this experiment are avoiding the additional risk of level collisions [135] inherent to our fixed coupling architecture [21] incurred by using the additional ancilla qubits required for X -type stabilizer measurements. This has been overcome in other experiments using different transmon frequency arrangements [38]. Despite this, the inherent crosstalk of fixed couplings will ultimately become unbearable. Instead, this issue can be made simpler by using tunable coupling architectures [10, 136–139].

6.2 Looking into the future

Quantum error correction offers a pathway to realize qubits with arbitrarily low errors [43]. It does so at the cost of high redundancy. Given the evolving landscape of quantum technologies, the feasibility of this approach is something worth reassessing. Adopting QEC is sensible if we cannot achieve a physical qubit system with inherently low errors and if the scaling required for such high redundancy is manageable. At this point, it's hard to make a definitive statement on either of these aspects. On one hand, realizing a low-error physical qubit likely

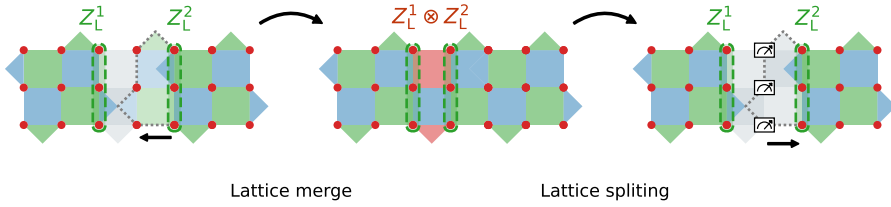


Figure 6.2: **Lattice surgery between ($d = 3$) surface code logical qubits.** Two distance $d = 3$ logical qubits, initially separated by a column of data qubits (red circles), are merged by measuring stabilizers (shaded tiles) along this boundary. This projects the state of the logical qubits along $Z_L \otimes Z_L$, which can be obtained from the the product of the stabilizer observables in red. The lattice is then split by measuring the column of data qubits which effectively removes them for

requires a major breakthrough in fundamental physics. A leading proposal for such a system involves realizing spatially separated Majorana fermions [144, 145], but it remains uncertain whether these states can be practically achieved. On the other hand, the scale necessary for useful, complex calculations with error-corrected qubits ($\sim 10^6$ qubits [143]) is far beyond current experimental achievements ($\sim 10^2$ qubits [146]). Some might even argue that the complexity of such scale may become unbearable. This problem is illustrated in Figure 6.1, where we plot the logical qubit error rate, ε_L , for a surface code as a function of the code distance, d , and the corresponding number of physical qubits, N , per logical qubit. Provided the physical qubit error is below the threshold of the code, $\varepsilon_L \propto \Lambda^{-(d+1)/2}$ [43, 44]. Even with an optimistic error suppression factor of $\Lambda = 3$ (similar to those achieved for current repetition codes [39]), we are looking at $N \gtrsim 4 \times 10^3$ for large-scale useful applications [62, 140–143]. Despite these challenges, I take comfort in the more tangible path towards realizing an error-corrected logical qubit. We have a clearer understanding of the technical requirements necessary to achieve it. As for which platform is best suited to achieve large-scale QEC, currently, superconducting qubits [38, 39, 146], trapped ions [91, 114], and neutral atom arrays [122, 123] seem to be enabling the largest scale demonstrations to date, positioning them as leading contenders. However, considering the current disparity between the projections in Figure 6.1 and actual experiments, it would be premature and shortsighted to dismiss any other existing quantum technologies [60, 61, 118, 121, 147, 148].

The conclusion of this discussion, which holds for all qubit technologies at the moment, is that the path forward in QEC will require both larger scales and lower physical error rates. Larger scale is essential to increase redundancy, which in turn enables lower logical error rates.

In this regard, advancements in 3D integration techniques [149–153] are critical for reducing qubit pitch on devices. Additionally, modular architectures [139, 154–156] offer promising solutions to overcome yield constraints. Simultaneously, achieving lower error rates is important for more effective error suppression. Research efforts in this direction include developing qubits with higher coherence times. Suppressing two-level system (TLS) defects [30–33, 134] will be particularly important to ensure uniformly low error rates across a device. Furthermore, explorations into new types of superconducting qubits, largely compatible with the circuit QED technology developed for transmons, have yielded encouraging results [157–162]. Another avenue that could significantly accelerate progress in QEC is the development of codes capable of encoding a greater number of logical qubits using the same number of physical qubits (denoted as N in Fig. 6.1). Quantum low-density parity-check (LDPC) codes, with higher encoding rates than the surface code [62], have recently seen a surge in research activity [163–165]. These codes typically achieve higher encoding rates by introducing longer-range connectivity between qubits, offering a promising direction for enhancing the efficiency of QEC codes thus alleviating the overhead required on scaling hardware.

While demonstrations of QEC using superconducting qubits have predominantly focused on memory experiments [38, 39, 90], the ultimate goal of an error-corrected quantum computer is to not only store but also process logical information [166]. As the scale of these systems increases, a natural next milestone is the demonstration of error-corrected logical operations [167]. Such advancements have already been achieved in systems using trapped ions [114] and neutral atoms [123]. The inherent all-to-all connectivity of these platforms facilitates transversal two-qubit gates between logical qubits. Conversely, in superconducting qubit systems, which typically feature nearest-neighbor connectivity, the most popular method for executing two-qubit logical gates is through lattice surgery [168]. This technique consists of merging and splitting the boundaries of logical qubits, as depicted in Figure 6.2. This allows performing $Z_L \otimes Z_L$ projective measurements on the logical qubits which can be used to generate entanglement [54].

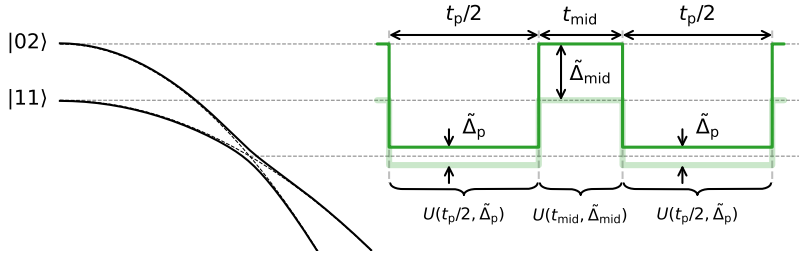


Figure A.1: **Trajectory of states during two-qubit gate.** Energy level diagram of bare states $|11\rangle$ and $|20\rangle$ of the two transmon system a function of flux (left) and their corresponding trajectory in time during two qubit gate. This trajectory can be split into three different parts each implementing a unitary transformation $U(t, \tilde{\Delta})$ that is a function of the level detuning $\tilde{\Delta}$ and duration t .

A.1 Analytical model of SNZ landscape.

Here I will drescribe the analytical model of the SNZ landscape. Consider the dynamics in the manifold $\{|11\rangle, |02\rangle\}$ given by the Hamiltonian,

$$H = \begin{pmatrix} -\tilde{\Delta}/2 & J_2 \\ J_2 & \tilde{\Delta}/2 \end{pmatrix}, \quad (\text{A.1})$$

where $\tilde{\Delta}$ and J_2 is the detuning and coupling between the two levels, respectively. To solve for the time evolution of H , it is usefull to diagonalize it using, $H = SMS^{-1}$, such that,

$$S = \begin{pmatrix} \cos(\theta/2) & -\sin(\theta/2) \\ \sin(\theta/2) & \cos(\theta/2) \end{pmatrix}, M = \begin{pmatrix} -\frac{\Omega}{2} & 0 \\ 0 & \frac{\Omega}{2} \end{pmatrix} \text{ and } S^{-1} = S^T, \quad (\text{A.2})$$

where, $\Omega = \sqrt{\tilde{\Delta}^2 + 4J_2^2}$ and $\tan(\theta) = 2J_2/\tilde{\Delta}$. The time evolution, $|\psi(t)\rangle$, of an initial state $|\psi_0\rangle$ is described by the unitary, U , such that $|\psi(t)\rangle = U|\psi_0\rangle$ where,

$$U(t, \tilde{\Delta}) = S \exp\left(+i\frac{\Omega}{2}\sigma_Z t\right) S^{-1}. \quad (\text{A.3})$$

To find the evolution of the state an initial state $|11\rangle$ a after the SNZ gate, we consider three distinct phases of the gate as shown in Figure A.1. Each phase will ideally implement a unitary $U(t, \tilde{\Delta})$ which will depend on the instant detuning $\tilde{\Delta}$ between levels and its the duration, t . Therefore, the state after the gate as function of the gate parameters will be given by,

$$|\psi\rangle(t_{\text{mid}}, \tilde{\Delta}_p, t_p, \tilde{\Delta}_{\text{mid}}, J_2) = U(t_p/2, \tilde{\Delta}_p)U(t_{\text{mid}}, \tilde{\Delta}_{\text{mid}})U(t_p/2, \tilde{\Delta}_p)|11\rangle. \quad (\text{A.4})$$

One can further simplify the resulting expression by using the approximation,

$$U(t_{\text{mid}}, \tilde{\Delta}_{\text{mid}}) \approx \exp\left(+i\frac{\tilde{\Delta}_{\text{mid}}}{2}\sigma_Z t_{\text{mid}}\right), \quad (\text{A.5})$$

which holds if the detuning at the sweetspot is much higher than the coupling between levels, i.e., $\tilde{\Delta}_{\text{mid}} \gg J_2$. In order to account for the effect of pulse distortion during the middle part of the gate, we introduce an additional parameter, $t_{\text{mid}}^{\text{dist}}$. This parameter is used as an offset such that $\tilde{t}_{\text{mid}} = t_{\text{mid}} - t_{\text{mid}}^{\text{dist}}$. Additionally, since in experiment one usually has direct control over the detuning of the highest frequency qubit, it is most convenient to express, $\tilde{\Delta}_{\text{p}} = \Delta_{\text{H}} - \Delta_0$. Here Δ_{H} and Δ_0 are the high-frequency qubit detuning and the detuning of the interaction, respectively. Equation A.4 turns into,

$$|\psi\rangle(t_{\text{mid}}, \Delta_{\text{H}}, \vec{\lambda}) = U(t_{\text{p}}/2, \tilde{\Delta}_{\text{p}}) e^{i \frac{\tilde{\Delta}_{\text{mid}}}{2} \sigma_z \tilde{t}_{\text{mid}}} U(t_{\text{p}}/2, \tilde{\Delta}_{\text{p}}) |11\rangle, \quad (\text{A.6})$$

where, $\vec{\lambda} = (\Delta_0, t_{\text{p}}, \tilde{\Delta}_{\text{mid}}, J_2, t_{\text{mid}}^{\text{dist}})$, is the vector of parameters that describe the gate. These will be used as free parameters in the fitting landscapes depicted in Figure 2.8. To do this, we must compute the average leakage of the gate,

$$L_1 = \frac{|\langle 02 | \psi \rangle|^2}{4}. \quad (\text{A.7})$$

To compute the conditional phase, we must calculate the phase acquired by the state $|11\rangle$ and subtract the single-qubit phase acquired by levels $|01\rangle$ and $|10\rangle$. This is done using:

$$\phi_{\text{c}} = \text{Arg}(\langle 11 | \psi \rangle) - \frac{\Delta_{\text{H}} t_{\text{p}}}{2} - \frac{\tilde{\Delta}_{\text{mid}} t_{\text{mid}}}{2}. \quad (\text{A.8})$$

where the last two terms correspond to the single-qubit phase acquired by $|01\rangle$ and $|10\rangle$ in our working frame of reference.

ACKNOWLEDGEMENTS

My time in Delft has been immensely fulfilling and joyful. Not only is Delft a great place to live and conduct cutting-edge research, but it is the people I met along the way who made it feel like a second home to me. First and foremost, I would like to thank my supervisor, **Leo**, for your guidance, support, and encouragement. You've taught me that determination is one of the most important traits of a researcher. I have learned a lot from you over the years, and I am very grateful to have had the opportunity to work together. The same goes for my co-promotor, **Barbara**. Thank you for your patience while I tried finding common ground between experiment and theory. The things I have learned within your group inspired much of my work, and I am grateful for all the support it has given me.

My journey in Delft started in the **Tittel** lab with **Gustavo**, **Anta**, and **Mohsen**. I was fortunate enough to have you as my first colleagues at TU Delft and even more fortunate to call you friends today. I will miss you all a lot, along with **Monique**, **Cecilia**, **Francisco**, **Maryieh** and **Mohammed**. QFA will never be the same without you. For all the joy one gets from doing research, it's hard maintaining sanity during times of stress. Fortunately, I could always count on my football teammates in the Quantum Warriors: **Gustavo**, **Francesco**, **Chris**, **Michael**, **Thijs**, **Maarten**, **Stefano**, **Alberto**, **Valentin**, **Bas**, **Sjoerd**, **Emmanuelle**, **Fabrizio**, **Maroto**, and **Vincent**, to help me blow off steam. It has been an honor playing by your side. I had a lot of fun during our matches and post-match beers. I trust your best season is yet to come. Halfway through my PhD, I picked up tennis. This sport feels like a steep (and slippery) learning curve that never seems to level out. To my tennis mentors, **Sjoerd** and **Francisco**: thank you for your patience. I can only hope to play as well as you by the time I'm forty (and my body allows it). Research hurdles appear in many forms, sometimes way too often. As scientists and engineers, we do the best we can, but sometimes your experimental setup just doesn't want to cooperate. To the electronics wizards, **Raymond** and **Raymond**: I don't know where I'd be without you. Needless to say, I have learned a lot from you. Also, a big thank you to **Jenny**, **Marja**, and **Chantele** for all the backstage support and for ensuring everything keeps running smoothly.

The mission-driven approach of Qutech to research in quantum technologies results in large combined efforts bringing together many fields of physics and engineering. The wide breadth of expertise required in these projects can only be fulfilled by effective collaborations. **Barbara**, **Francesco**, and **Boris**, I think our work is a great demonstration of this, and I would like to thank you for that. **Francesco**, thank you for coming up with an LRU. Let me know if there is anything else you are cooking. **Boris**, thank you for always being there to help.

You made quantum error correction understandable for me. I have really enjoyed our work together, and I am excited to continue collaborating in the future.

A special thank you to our friends across the hallway: **Figen, Taryn, Eugene, Martijn, Siddhart, and Christian**. I was really excited when I first heard you were moving in. We spent some great times together in and outside the lab. I wish you great success in your research, and I look forward to following your progress.

In research, as in life, one can accomplish much more as a team. To the entire DiCarlo Lab family, which I was proud to be a part of: **Hany, Miguel, Thijs, Santi, Ruggero, Sean, Tim, Bart, Nandini, Rebecca, Alessandro, Ramiro, Marc, Matvey, Martijn, Christos, Niels, Adriaan, Dickel, Wouter, Victor, Berend, Tumi, Pim, Hires, Joost, Olexyi, Yuejie, Marios, Viswanath, Kishore and Leo**. **Sean**, your automation skills are unmatched (as are your dancing moves). We could not have asked for a better successor in Surface-17. Thank you and **Zalyna** for making Delft such a *gezellig* place for us. **Tim**, I am proud you were my first master's student and was really happy you decided to stick around for the long ride. It's great to see you evolve as a researcher and to have you as a friend. I have no doubt you'll accomplish great things in your career. **Bart**, thank you for your friendship. I really enjoyed the time we spent together. I hope you find happiness in whatever you decide to do next in life. To the first person I met in the group, **Thijs**: so many things to thank you for. You were an awesome colleague, mentor, beer pal, striker and of course - most of all - friend. I will miss you a lot. To my good friends and paranympths **Santi** and **Ruggero**: it's not often in life that I connect with someone the way I have with you. We shared the happiest memories I have of Delft. I think the best way to express this is by sharing one of those memories. It's Friday afternoon and the clock hits 5 pm. **Santi**'s drinking bell rings in an attempt to shame those who dare to work past that time on a Friday. We all head to TPKV, picking up lost folks along the way, to drink a couple of pitchers before heading to the Beestenmarkt to grab dinner. There is a pointless attempt to convince me that McDonald's is not the best restaurant in Delft. After dinner, together with **Thijs**, we head to 't **Proeflokaal**. Tensions are high at this point. A dart match is about to start. But not without first matching tension and blood alcohol levels. The scoreboard is filling up, leaving only the Bulls. **Srijit** joins the match at this point, later beating all of us. A humbling experience and **Santi** is taking it the hardest. In the aftermath of the events, we find ourselves draining the remaining beer left in the draft pipes. I get a triple, **Santi** gets a hazy IPA and **Ruggero** asks for a weird chocolate "beer". It's now 2 o'clock and the nieuwe kerk bell rings. Sadly, we must leave 't **Proeflokaal**. Our options are scarce at this point, but we carry the party to **Oude Jan**. Upon getting there, we find **Bart** in the crowd. It's going to be a long night... After this point, my recollection of events is not faithful anymore. I think someone might have even lost his shirt at some point. Anyways, I'll always remember these and other moments we spent together. I'll miss you a lot and wish you both (and **Sara**) tons of success in your life.

To **Hany**, my PhD buddy. We have been through a lot together. Thank you for being there

throughout the good and bad moments. You were a great colleague and an even better friend. I hope our paths cross again in life. Until then, keep enjoying the best that life has to offer, little **Zayn** and **Jawahir**.

Ao **Miguel** e à **Leonor**. **Miguel**, foste um colega incansável e um grande amigo durante este percurso. Modéstia à parte, acho que fazemos uma grande equipa juntos. Espero estar à altura de te retribuir o favor nesta nova aventura. Estou muito feliz que possamos estar todos juntos outra vez e espero que os nossos churrascos continuem a correr o mundo.

To my friends **Xico** and **Magdalena**. **Xico**, existe uma Holanda antes e depois de ti. Há amigos que ficam para sempre e espero que seja sempre assim connosco. You guys made the dutch weather look a lot less grey and we'll miss you a lot. Our door is always open for you and we hope you can come visit soon!

À minha família: Mãe, Pai, Cláudia e Cláudia. Esta foi mais uma etapa à qual vos devo muito e em que a parte mais difícil foram as saudades de voltar a casa. Obrigado por acreditarem em mim e sempre me encorajarem a ir mais longe.

E finalmente, à **Andreia**. Obrigado por seres a minha companheira nesta aventura. Não queria (nem conseguiria) fazer isto sem ti. Vou-me lembrar de Delft como uma altura muito feliz das nossas vidas e estou ansioso por tudo o que ainda vamos viver juntos.

Jorge MIGUEL FERREIRA MARQUES

27.12.1995 Born in Porto, Portugal.

Education

2013-2016 Bachelor of science in physics
Faculdade de Ciências da Universidade do Porto, Porto

2016-2019 Master of science in engineering physics
Insituto Superior Técnico, Lisboa
Master end project:
Thesis: Measurement-device-independent quantum key distribution
Advisor: Prof. dr. Y. Omar

2019-2024 Ph.D. in experimental Physics
Technische Universiteit Delft
Thesis: Quantum error correction with superconducting qubits
Advisors: Prof. dr. L. DiCarlo and Prof. dr. B. M. Terhal.

LIST OF PUBLICATIONS

7. H. Ali, **J. F. Marques**, O. Crawford, J. Majaniemi, M. Serra-Peralta, D. Byfield, B. M. Varbanov, B. M. Terhal, E. T. Campbell, and L. DiCarlo, *Reducing the error rate of a superconducting logical qubit using analog readout information*, Phys. Rev. Appl., (in press, 2024).
6. M. S. Moreira, G. G. Guerreschi, W. Vlothuizen, **J. F. Marques**, J. van Straten, S. P. Premaratne, X. Zou, H. Ali, N. Muthusubramanian, C. Zachariadis, J van Someren, M. Beekman, N. Haider, A. Bruno, C. G. Almudever, A. Y. Matsuura, L. DiCarlo, *Realization of a quantum neural network using repeat-until-success circuits in a superconducting quantum processor*, npj Quantum Information **9**, 118 (2023).
5. **J. F. Marques**, H. Ali, B. M. Varbanov, M. Finkel, H. M. Veen, S. L. M. van der Meer, S. Vallés-Sanclemente, N. Muthusubramanian, M. Beekman, N. Haider, B. M. Terhal, L. DiCarlo, *All-microwave leakage reduction units for quantum error correction with superconducting transmon qubits*, Physical Review Letters **130**, 250602 (2023).
4. S. Vallés-Sanclemente, S. L. M. van der Meer, M. Finkel, N. Muthusubramanian, M. Beekman, H. Ali, **J. F. Marques**, C. Zachariadis, H. M. Veen, T. Stavenga, N. Haider, L. DiCarlo, *Post-fabrication frequency trimming of coplanar-waveguide resonators in circuit QED quantum processors*, Applied Physics Letters **123**, 034004 (2023).
3. **J. F. Marques**, B. M. Varbanov, M. S. Moreira, H. Ali, N. Muthusubramanian, C. Zachariadis, F. Battistel, M. Beekman, N. Haider, W. Vlothuizen, A. Bruno, B. M. Terhal, L. DiCarlo, *Logical-qubit operations in an error-detecting surface code*, Nature Physics **18**, 80–86 (2022).
2. R. C. Berrevoets, T. Middelburg, R. F. L. Vermeulen, L. D. Chiesa, F. Broggi, S. Piciaccia, R. Pluis, P. Umesh, **J. F. Marques**, W. Tittel, J. A. Slater, *Deployed measurement-device independent quantum key distribution and Bell-state measurements coexisting with standard internet data and networking equipment*, Communications Physics **5**, 186 (2022).
1. V. Negîrneac, H. Ali, N. Muthusubramanian, F. Battistel, R. Sagastizabal, M. S. Moreira, **J. F. Marques**, W. Vlothuizen, M. Beekman, C. Zachariadis, N. Haider, A. Bruno, L. DiCarlo, *High-fidelity controlled-Z gate with maximal intermediate leakage operating at the speed limit in a superconducting quantum processor*, Physical Review Letters **126**, 220502 (2021).

REFERENCES

- [1] A. BLAIS, J. GAMBETTA, A. WALLRAFF, D. I. SCHUSTER, S. M. GIRVIN, M. H. DEVORET, and R. J. SCHOELKOPF. Quantum-information processing with circuit quantum electrodynamics. *Phys. Rev. A*, **75**, 032329, 2007.
- [2] A. BLAIS, A. L. GRIMSMO, S. M. GIRVIN, and A. WALLRAFF. *Circuit quantum electrodynamics*. *Rev. Mod. Phys.*, **93**, 025005, 2021.
- [3] J. KOCH, T. M. YU, J. GAMBETTA, A. A. HOUCK, D. I. SCHUSTER, J. MAJER, A. BLAIS, M. H. DEVORET, S. M. GIRVIN, and R. J. SCHOELKOPF. *Charge-insensitive qubit design derived from the Cooper pair box*. *Phys. Rev. A*, **76**, 042319, 2007.
- [4] B. D. JOSEPHSON. *Possible new effects in superconductive tunnelling*. *Phys. Lett.*, **1** (7), 251–253, 1962.
- [5] M. TINKHAM. *Introduction to Superconductivity*. McGraw-Hill, New York, 2nd edition, 1996.
- [6] U. VOOL and M. DEVORET. *Introduction to quantum electromagnetic circuits*. *International Journal of Circuit Theory and Applications*, **45** (7), 897–934, 2017.
- [7] L. DICARLO, J. M. CHOW, J. M. GAMBETTA, L. S. BISHOP, B. R. JOHNSON, D. I. SCHUSTER, J. MAJER, A. BLAIS, L. FRUNZIO, S. M. GIRVIN, and R. J. SCHOELKOPF. *Demonstration of two-qubit algorithms with a superconducting quantum processor*. *Nature*, **460**, 240, 2009.
- [8] R. BARENDTS, J. KELLY, A. MEGRANT, A. VEITIA, D. SANK, E. JEFFREY, T. C. WHITE, J. MUTUS, A. G. FOWLER, B. CAMPBELL, Y. CHEN, Z. CHEN, B. CHIARO, *et al.* *Superconducting quantum circuits at the surface code threshold for fault tolerance*. *Nature*, **508** (7497), 500, 2014.
- [9] M. A. ROL, F. BATTISTEL, F. K. MALINOWSKI, C. C. BULTINK, B. M. TARASINSKI, R. VOLLMER, N. HAIDER, N. MUTHUSUBRAMANIAN, A. BRUNO, B. M. TERHAL, and L. DICARLO. *Fast, high-fidelity conditional-phase gate exploiting leakage interference in weakly anharmonic superconducting qubits*. *Phys. Rev. Lett.*, **123**, 120502, 2019.
- [10] Y. SUNG, L. DING, J. BRAUMÜLLER, A. VEPSÄLÄINEN, B. KANNAN, M. KJAERGAARD, A. GREENE, G. O. SAMACH, C. MCNALLY, D. KIM, A. MELVILLE, B. M. NIEDZIELSKI, M. E. SCHWARTZ, *et al.* *Realization of high-fidelity cz and zz-free iswap gates with a tunable coupler*. *Phys. Rev. X*, **11**, 021058, 2021.
- [11] V. NEGÎRNEAC, H. ALI, N. MUTHUSUBRAMANIAN, F. BATTISTEL, R. SAGASTIZABAL, M. S. MOREIRA, J. F. MARQUES, W. J. VLOTHUIZEN, M. BEEKMAN, C. ZACHARIADIS, N. HAIDER, A. BRUNO, and L. DICARLO. *High-fidelity controlled-z gate with maximal intermediate leakage operating at the speed limit in a superconducting quantum processor*. *Phys. Rev. Lett.*, **126**, 220502, 2021.
- [12] D. P. DIVINCENZO. *The physical implementation of quantum computation*. *Fortschritte der Physik*, **48**, 771, 2000.
- [13] P. KRANTZ, M. KJAERGAARD, F. YAN, T. P. ORLANDO, S. GUSTAVSSON, and W. D. OLIVER. *A quantum engineer's guide to superconducting qubits*. *Applied Physics Reviews*, **6** (2), 2019.
- [14] D. M. POZAR. *Microwave engineering*. John Wiley & Sons, Hoboken, 4 edition, 2005.

- [15] J. M. MARTINIS and M. R. GELLER. *Fast adiabatic qubit gates using only σ_z control*. Phys. Rev. A, **90**, 022307, 2014.
- [16] R. BARENDT, C. M. QUINTANA, A. G. PETUKHOV, Y. CHEN, D. KAFRI, K. KECHEDZHI, R. COLLINS, O. NAAMAN, S. BOIXO, F. ARUTE, K. ARYA, D. BUELL, B. BURKETT, *et al.* *Diabatic gates for frequency-tunable superconducting qubits*. Phys. Rev. Lett., **123**, 210501, 2019.
- [17] N. LACROIX, C. HELLINGS, C. K. ANDERSEN, A. DI PAOLO, A. REMM, S. LAZAR, S. KRINNER, G. J. NORRIS, M. GABUREAC, J. HEINSOO, A. BLAIS, C. EICHLER, and A. WALLRAFF. *Improving the performance of deep quantum optimization algorithms with continuous gate sets*. PRX Quantum, **1**, 020304, 2020.
- [18] D. SANK, Z. CHEN, M. KHEZRI, J. KELLY, R. BARENDT, B. CAMPBELL, Y. CHEN, B. CHIARO, A. DUNSWORTH, A. FOWLER, *et al.* *Measurement-induced state transitions in a superconducting qubit: Beyond the rotating wave approximation*. Physical review letters, **117** (19), 190503, 2016.
- [19] M. KHEZRI, A. OPREMCÁK, Z. CHEN, K. C. MIAO, M. MCEWEN, A. BENGTSSON, T. WHITE, O. NAAMAN, D. SANK, A. N. KOROTKOV, *et al.* *Measurement-induced state transitions in a superconducting qubit: Within the rotating-wave approximation*. Physical Review Applied, **20** (5), 054008, 2023.
- [20] M. F. DUMAS, B. GROLEAU-PARÉ, A. McDONALD, M. H. MUÑOZ-ARIAS, C. LLEDÓ, B. D’ANJOU, and A. BLAIS. *Unified picture of measurement-induced ionization in the transmon*. arXiv preprint arXiv:2402.06615, 2024.
- [21] R. VERSLUIS, S. POLETO, N. KHAMMASSI, B. TARASINSKI, N. HAIDER, D. J. MICHALAK, A. BRUNO, K. BERTELS, and L. DICARLO. *Scalable quantum circuit and control for a superconducting surface code*. Phys. Rev. Appl., **8**, 034021, 2017.
- [22] M. REED. *Entanglement and quantum error correction with superconducting qubits*. PhD Dissertation, Yale University, 2013.
- [23] F. MOTZOI, J. M. GAMBETTA, P. REBENTROST, and F. K. WILHELM. *Simple pulses for elimination of leakage in weakly nonlinear qubits*. Phys. Rev. Lett., **103**, 110501, 2009.
- [24] E. KNILL, D. LEIBFRIED, R. REICHLE, J. BRITTON, R. B. BLAKESTAD, J. D. JOST, C. LANGER, R. OZERI, S. SEIDELIN, and D. J. WINELAND. *Randomized benchmarking of quantum gates*. Phys. Rev. A, **77**, 012307, 2008.
- [25] E. MAGESAN, J. M. GAMBETTA, and J. EMERSON. *Scalable and robust randomized benchmarking of quantum processes*. Phys. Rev. Lett., **106**, 180504, 2011.
- [26] E. MAGESAN, J. M. GAMBETTA, and J. EMERSON. *Characterizing quantum gates via randomized benchmarking*. Phys. Rev. A, **85**, 042311, 2012.
- [27] E. MAGESAN, J. M. GAMBETTA, B. R. JOHNSON, C. A. RYAN, J. M. CHOW, S. T. MERKEL, M. P. DA SILVA, G. A. KEEFE, M. B. ROTHWELL, T. A. OHKI, M. B. KETCHEN, and M. STEFFEN. *Efficient measurement of quantum gate error by interleaved randomized benchmarking*. Phys. Rev. Lett., **109**, 080505, 2012.
- [28] M. A. ROL, L. CIORCIARO, F. K. MALINOWSKI, B. M. TARASINSKI, R. E. SAGASTIZABAL, C. C. BULTINK, Y. SALATHE, N. HAANDBAEK, J. SEDIVY, and L. DICARLO. *Time-domain characterization and correction of on-chip distortion of control pulses in a quantum processor*. App. Phys. Lett., **116** (5), 054001, 2020.
- [29] W. A. PHILLIPS. *Two-level states in glasses*. Rep. Prog. Phys., **50** (12), 1657, 1987.

- [30] J. LISENFELD, A. BILMES, A. MEGRANT, R. BARENDs, J. KELLY, P. KLIMOV, G. WEISS, J. M. MARTINIS, and A. V. USTINOV. *Electric field spectroscopy of material defects in transmon qubits*. npj Quantum Information, **5** (1), 2019.
- [31] A. BILMES, A. MEGRANT, P. KLIMOV, G. WEISS, J. M. MARTINIS, A. V. USTINOV, and J. LISENFELD. Resolving the positions of defects in superconducting quantum bits. Scientific reports, **10** (1), 3090, 2020.
- [32] A. BILMES, S. VOLOSHENIUK, A. V. USTINOV, and J. LISENFELD. Probing defect densities at the edges and inside josephson junctions of superconducting qubits. npj Quantum Information, **8** (1), 24, 2022.
- [33] J. LISENFELD, A. BILMES, and A. V. USTINOV. Enhancing the coherence of superconducting quantum bits with electric fields. npj Quantum Information, **9** (1), 8, 2023.
- [34] J. JÄCKLE. On the ultrasonic attenuation in glasses at low temperatures. Zeitschrift für Physik A Hadrons and nuclei, **257** (3), 212–223, 1972.
- [35] M. CHEN, J. C. OWENS, H. PUTTERMAN, M. SCHÄFER, and O. PAINTER. Phonon engineering of atomic-scale defects in superconducting quantum circuits. arXiv preprint arXiv:2310.03929, 2023.
- [36] A. P. PLACE, L. V. RODGERS, P. MUNDADA, B. M. SMITHAM, M. FITZPATRICK, Z. LENG, A. PREMKUMAR, J. BRYON, A. VRAJITOAREA, S. SUSSMAN, *et al.* New material platform for superconducting transmon qubits with coherence times exceeding 0.3 milliseconds. Nature communications, **12** (1), 1779, 2021.
- [37] A. BRUNO, G. DE LANGE, S. ASAAD, K. L. VAN DER ENDEN, N. K. LANGFORD, and L. DICARLO. *Reducing intrinsic loss in superconducting resonators by surface treatment and deep etching of silicon substrates*. Appl. Phys. Lett., **106**, 182 601, 2015.
- [38] S. KRINNER, N. LACROIX, A. REMM, D. P. AGUSTIN, E. GENOIS, C. LEROUX, C. HELLINGS, S. LAZAR, F. SWIADEK, J. HERRMANN, G. J. NORRIS, C. K. ANDERSEN, M. MÜLLER, *et al.* *Realizing repeated quantum error correction in a distance-three surface code*. Nature, **605** (7911), 669–674, 2022.
- [39] R. ACHARYA, I. ALEINER, R. ALLEN, T. ANDERSEN, M. ANSMANN, F. ARUTE, *et al.* Suppressing quantum errors by scaling a surface code logical qubit. Nature, **614** (7949), 676–681, 2023.
- [40] T. THORBECK, Z. XIAO, A. KAMAL, and L. C. G. GOVIA. Readout-induced suppression and enhancement of superconducting qubit lifetimes. arXiv preprint arXiv:2305.10508, 2023.
- [41] J. HEINSOO, C. K. ANDERSEN, A. REMM, S. KRINNER, T. WALTER, Y. SALATHÉ, S. GASPARINETTI, J.-C. BESSE, A. POTOČNIK, A. WALLRAFF, and C. EICHLER. *Rapid high-fidelity multiplexed readout of superconducting qubits*. Phys. Rev. Appl., **10**, 034 040, 2018.
- [42] J. PRESKILL. *Quantum Computing in the NISQ era and beyond*. Quantum, **2**, 79, 2018.
- [43] B. M. TERHAL. *Quantum error correction for quantum memories*. Rev. Mod. Phys., **87**, 307–346, 2015.
- [44] J. M. MARTINIS. *Qubit metrology for building a fault-tolerant quantum computer*. npj Quantum Inf., **1**, 15 005, 2015.
- [45] J. KELLY, R. BARENDs, A. FOWLER, A. MEGRANT, E. JEFFREY, T. WHITE, D. SANK, J. MUTUS, B. CAMPBELL, Y. CHEN, *et al.* *State preservation by repetitive error detection in a superconducting quantum circuit*. Nature, **519** (7541), 66–69, 2015.
- [46] Z. CHEN, K. J. SATZINGER, J. ATALAYA, A. N. KOROTKOV, A. DUNSWORTH, D. SANK, C. QUINTANA, M. MCEWEN, R. BARENDs, P. V. KLIMOV, S. HONG, C. JONES, A. PETUKHOV, *et al.* *Exponential suppression of bit or phase errors with cyclic error correction*. Nature, **595** (7867), 383–387, 2021.

- [47] R. LESCANNE, M. VILLIERS, T. PERONNIN, A. SARLETTE, M. DELBECQ, B. HUARD, T. KONTOS, M. MIRRAHIMI, and Z. LEGHTAS. *Exponential suppression of bit-flips in a qubit encoded in an oscillator*. *Nature Physics*, **16** (5), 509–513, 2020.
- [48] A. GRIMM, N. E. FRATTINI, S. PURI, S. O. MUNDHADA, S. TOUZARD, M. MIRRAHIMI, S. M. GIRVIN, S. SHANKAR, and M. H. DEVORET. *Stabilization and operation of a kerr-cat qubit*. *Nature*, **584** (7820), 205–209, 2020.
- [49] D. RISTÈ, S. POLETTI, M. Z. HUANG, A. BRUNO, V. VESTERINEN, O. P. SAIRA, and L. DICARLO. *Detecting bit-flip errors in a logical qubit using stabilizer measurements*. *Nat. Comm.*, **6**, 6983, 2015.
- [50] J. CRAMER, N. KALB, M. A. ROL, B. HENSEN, M. MARKHAM, D. J. TWITCHEN, R. HANSON, and T. H. TAMINIAU. *Repeated quantum error correction on a continuously encoded qubit by real-time feedback*. *Nat. Comm.*, **5**, 11 526, 2016.
- [51] D. RISTÈ, L. C. G. GOVIA, B. DONOVAN, S. D. FALLEK, W. D. KALFUS, M. BRINK, N. T. BRONN, and T. A. OHKI. *Real-time processing of stabilizer measurements in a bit-flip code*. *npj Quantum Information*, **6** (1), 71, 2020.
- [52] D. NIGG, M. MÜLLER, E. A. MARTINEZ, P. SCHINDLER, M. HENNRICH, T. MONZ, M. A. MARTIN-DELGADO, and R. BLATT. *Quantum computations on a topologically encoded qubit*. *Science*, **345** (6194), 302–305, 2014.
- [53] L. EGAN, D. M. DEBROY, C. NOEL, A. RISINGER, D. ZHU, D. BISWAS, M. NEWMAN, M. LI, K. R. BROWN, M. CETINA, *et al.* *Fault-tolerant control of an error-corrected qubit*. *Nature*, **598** (7880), 281–286, 2021.
- [54] A. ERHARD, H. P. NAUTRUP, M. METH, L. POSTLER, R. STRICKER, M. STADLER, V. NEGNEVITSKY, M. RINGBAUER, P. SCHINDLER, H. J. BRIEGEL, *et al.* *Entangling logical qubits with lattice surgery*. *Nature*, **589** (7841), 220–224, 2021.
- [55] V. NEGNEVITSKY, M. MARINELLI, K. K. MEHTA, H.-Y. LO, C. FLÜHMANN, and J. P. HOME. *Repeated multi-qubit readout and feedback with a mixed-species trapped-ion register*. *Nature*, **563** (7732), 527–531, 2018.
- [56] A. BLAIS, R.-S. HUANG, A. WALLRAFF, S. M. GIRVIN, and R. J. SCHOELKOPF. *Cavity quantum electrodynamics for superconducting electrical circuits: An architecture for quantum computation*. *Phys. Rev. A*, **69**, 062 320, 2004.
- [57] C. K. ANDERSEN, A. REMM, S. LAZAR, S. KRINNER, J. HEINSOO, J.-C. BESSE, M. GABUREAC, A. WALLRAFF, and C. EICHLER. *Entanglement stabilization using ancilla-based parity detection and real-time feedback in superconducting circuits*. *npj Quantum Information*, **5** (69), 1–7, 2019.
- [58] C. C. BULTINK, T. E. O'BRIEN, R. VOLLMER, N. MUTHUSUBRAMANIAN, M. W. BEEKMAN, M. A. ROL, X. FU, B. TARASINSKI, V. OSTROUKH, B. VARBANOV, A. BRUNO, and L. DICARLO. *Protecting quantum entanglement from leakage and qubit errors via repetitive parity measurements*. *Science Advances*, **6** (12), 2020.
- [59] N. OFEK, A. PETRENKO, R. HEERES, P. REINHOLD, Z. LEGHTAS, B. VLASTAKIS, Y. LIU, L. FRUNZIO, S. M. GIRVIN, L. JIANG, M. MIRRAHIMI, M. H. DEVORET, and R. J. SCHOELKOPF. *Extending the lifetime of a quantum bit with error correction in superconducting circuits*. *Nature*, **536**, 441, 2016.
- [60] *Quantum error correction and universal gate set operation on a binomial bosonic logical qubit*. *Nat. Phys.*, 2019.
- [61] P. CAMPAGNE-IBARCO, A. EICKBUSCH, S. TOUZARD, E. ZALYS-GELLER, N. E. FRATTINI, V. V. SIVAK, P. REINHOLD, S. PURI, S. SHANKAR, R. J. SCHOELKOPF, L. FRUNZIO, M. MIRRAHIMI, and M. H. DEVORET. *Quantum error correction of a qubit encoded in grid states of an oscillator*. *Nature*, **584** (7821), 368–372, 2020.

- [62] A. G. FOWLER, M. MARIANTONI, J. M. MARTINIS, and A. N. CLELAND. *Surface codes: Towards practical large-scale quantum computation*. Phys. Rev. A, **86**, 032324, 2012.
- [63] C. K. ANDERSEN, A. REMM, S. LAZAR, S. KRINNER, N. LACROIX, G. J. NORRIS, M. GABUREAC, C. EICHLER, and A. WALLRAFF. *Repeated quantum error detection in a surface code*. Nat. Phys., **16** (8), 875–880, 2020.
- [64] J. M. CHOW, J. M. GAMBETTA, A. D. C R COLES, S. T. MERKEL, J. A. SMOLIN, C. RIGETTI, S. POLETTI, G. A. KEEFE, M. B. ROTHWELL, J. R. ROZEN, M. B. KETCHEN, and M. STEFFEN. *Universal quantum gate set approaching fault-tolerant thresholds with superconducting qubits*. Phys. Rev. Lett., **109**, 060501, 2012.
- [65] Y. TOMITA and K. M. SVORE. *Low-distance surface codes under realistic quantum noise*. Phys. Rev. A, **90**, 062320, 2014.
- [66] O.-P. SAIRA, J. P. GROEN, J. CRAMER, M. MERETSKA, G. DE LANGE, and L. DICARLO. *Entanglement genesis by ancilla-based parity measurement in 2D circuit QED*. Phys. Rev. Lett., **112**, 070502, 2014.
- [67] M. TAKITA, A. D. C R COLES, E. MAGESAN, B. ABDO, M. BRINK, A. CROSS, J. M. CHOW, and J. M. GAMBETTA. *Demonstration of weight-four parity measurements in the surface code architecture*. arXiv:1605.01351, 2016.
- [68] P. ALIFERIS, D. GOTTESMAN, and J. PRESKILL. *Quantum accuracy threshold for concatenated distance-3 codes*. Quantum Inf. Comput., **6** (quant-ph/0504218), 97–165, 2005.
- [69] A. PAETZNICK and K. M. SVORE. *Repeat-until-success: non-deterministic decomposition of single-qubit unitaries*. Quantum Information & Computation, **14**, 1277–1301, 2014.
- [70] B. M. VARBANOV, F. BATTISTEL, B. M. TARASINSKI, V. P. OSTROUKH, T. E. O'BRIEN, L. DICARLO, and B. M. TERHAL. *Leakage detection for a transmon-based surface code*. npj Quantum Information, **6** (1), 102, 2020.
- [71] M. MCEWEN, D. KAFRI, Z. CHEN, J. ATALAYA, K. SATZINGER, C. QUINTANA, P. V. KLIMOV, D. SANK, C. GIDNEY, A. FOWLER, *et al.* *Removing leakage-induced correlated errors in superconducting quantum error correction*. Nature communications, **12** (1), 1–7, 2021.
- [72] F. BATTISTEL, B. VARBANOV, and B. TERHAL. *Hardware-efficient leakage-reduction scheme for quantum error correction with superconducting transmon qubits*. PRX Quantum, **2**, 030314, 2021.
- [73] K. J. SATZINGER, Y.-J. LIU, A. SMITH, C. KNAPP, M. NEWMAN, C. JONES, Z. CHEN, C. QUINTANA, X. MI, A. DUNSWORTH, C. GIDNEY, I. ALEINER, F. ARUTE, *et al.* *Realizing topologically ordered states on a quantum processor*. Science, **374** (6572), 1237–1241, 2021.
- [74] J. A. SCHREIER, A. A. HOUCK, J. KOCH, D. I. SCHUSTER, B. R. JOHNSON, J. M. CHOW, J. M. GAMBETTA, J. MAJER, L. FRUNZIO, M. H. DEVORET, S. M. GIRVIN, and R. J. SCHOELKOPF. *Suppressing charge noise decoherence in superconducting charge qubits*. Phys. Rev. B, **77**, 180502, 2008.
- [75] Qiskit: An open-source framework for quantum computing, 2019.
- [76] J. DE JONG. *Implementation of a fault-tolerant SWAP operation on the IBM 5-qubit device*. Master's thesis, Delft University of Technology, 2019.
- [77] C. C. BULTINK, B. TARASINSKI, N. HAANDBAEK, S. POLETTI, N. HAIDER, D. J. MICHALAK, A. BRUNO, and L. DICARLO. *General method for extracting the quantum efficiency of dispersive qubit readout in circuit qed*. Appl. Phys. Lett., **112**, 092601, 2018.
- [78] D. RIST , J. G. VAN LEEUWEN, H.-S. KU, K. W. LEHNERT, and L. DICARLO. *Initialization by measurement of a superconducting quantum bit circuit*. Phys. Rev. Lett., **109**, 050507, 2012.

- [79] T. WALTER, P. KURPIERS, S. GASPARINETTI, P. MAGNARD, A. POTOČNIK, Y. SALATHÉ, M. PECHAL, M. MONDAL, M. OPPLIGER, C. EICHLER, and A. WALLRAFF. *Rapid High-Fidelity Single-Shot Dispersive Readout of Superconducting Qubits*. Phys. Rev. Appl., **7** (5), 054 020, 2017.
- [80] R. SAGASTIZABAL, S. P. PREMARATNE, B. A. KLAVER, M. A. ROL, V. NEGÎRNEAC, M. MOREIRA, X. ZOU, S. JOHRI, N. MUTHUSUBRAMANIAN, M. BEEKMAN, C. ZACHARIADIS, V. P. OSTROUKH, N. HAIDER, *et al.* *Variational preparation of finite-temperature states on a quantum computer*. npj Quantum Inf., **7**, 130, 2021.
- [81] C. J. WOOD and J. M. GAMBETTA. *Quantification and characterization of leakage errors*. Phys. Rev. A, **97**, 032 306, 2018.
- [82] S. ASAAD, C. DICKEL, S. POLETO, A. BRUNO, N. K. LANGFORD, M. A. ROL, D. DEURLOO, and L. DICARLO. *Independent, extensible control of same-frequency superconducting qubits by selective broadcasting*. npj Quantum Inf., **2**, 16 029, 2016.
- [83] T. O'BRIEN, B. TARASINSKI, and L. DICARLO. *Density-matrix simulation of small surface codes under current and projected experimental noise*. npj Quantum Information, **3** (1), 39, 2017.
- [84] M. TAKITA, A. W. CROSS, A. D. CÔRCOLES, J. M. CHOW, and J. M. GAMBETTA. *Experimental demonstration of fault-tolerant state preparation with superconducting qubits*. Phys. Rev. Lett., **119**, 180 501, 2017.
- [85] P. ALIFERIS and B. M. TERHAL. *Fault-tolerant quantum computation for local leakage faults*. Quantum Info. Comput., **7** (1), 139–156, 2007.
- [86] A. G. FOWLER. *Coping with qubit leakage in topological codes*. Phys. Rev. A, **88**, 042 308, 2013.
- [87] J. GHOSH, A. G. FOWLER, J. M. MARTINIS, and M. R. GELLER. *Understanding the effects of leakage in superconducting quantum-error-detection circuits*. Phys. Rev. A, **88**, 062 329, 2013.
- [88] C. RYAN-ANDERSON, J. G. BOHNET, K. LEE, D. GRESH, A. HANKIN, J. P. GAEBLER, D. FRANCOIS, A. CHERNOGUZOV, D. LUCCHETTI, N. C. BROWN, T. M. GATTERMAN, S. K. HALIT, K. GILMORE, *et al.* *Realization of real-time fault-tolerant quantum error correction*. Phys. Rev. X, **11**, 041 058, 2021.
- [89] Y. ZHAO, Y. YE, H.-L. HUANG, Y. ZHANG, D. WU, H. GUAN, Q. ZHU, Z. WEI, T. HE, S. CAO, F. CHEN, T.-H. CHUNG, H. DENG, *et al.* *Realization of an error-correcting surface code with superconducting qubits*. Phys. Rev. Lett., **129**, 030 501, 2022.
- [90] N. SUNDARESAN, T. J. YODER, Y. KIM, M. LI, E. H. CHEN, G. HARPER, T. THORBECK, A. W. CROSS, A. D. CÔRCOLES, and M. TAKITA. *Demonstrating multi-round subsystem quantum error correction using matching and maximum likelihood decoders*. Nature Communications, **14** (1), 2852, 2023.
- [91] L. POSTLER, S. HEUßEN, I. POGORELOV, M. RISPLER, T. FELDKER, M. METH, C. D. MARCINIAK, R. STRICKER, M. RINGBAUER, R. BLATT, P. SCHINDLER, M. MÜLLER, and T. MONZ. *Demonstration of fault-tolerant universal quantum gate operations*. Nature, **605** (7911), 675–680, 2022.
- [92] P. MAGNARD, P. KURPIERS, B. ROYER, T. WALTER, J.-C. BESSE, S. GASPARINETTI, M. PECHAL, J. HEINSOO, S. STORZ, A. BLAIS, and A. WALLRAFF. *Fast and unconditional all-microwave reset of a superconducting qubit*. Phys. Rev. Lett., **121**, 060 502, 2018.
- [93] J. GHOSH and A. G. FOWLER. *Leakage-resilient approach to fault-tolerant quantum computing with superconducting elements*. Phys. Rev. A, **91**, 020 302(R), 2015.
- [94] M. SUCHARA, A. W. CROSS, and J. M. GAMBETTA. *Leakage suppression in the toric code*. Quantum Info. Comput., **15** (11-12), 997–1016, 2015.

- [95] M. McEWEN, D. BACON, and C. GIDNEY. *Relaxing hardware requirements for surface code circuits using time-dynamics*, 2023.
- [96] N. C. BROWN, M. NEWMAN, and K. R. BROWN. *Handling leakage with subsystem codes*. *New Journal of Physics*, **21** (7), 073055, 2019.
- [97] N. C. BROWN, A. CROSS, and K. R. BROWN. Critical faults of leakage errors on the surface code. In *2020 IEEE International Conference on Quantum Computing and Engineering (QCE)*, pages 286–294. 2020.
- [98] D. HAYES, D. STACK, B. BJORK, A. C. POTTER, C. H. BALDWIN, and R. P. STUTZ. *Eliminating leakage errors in hyperfine qubits*. *Phys. Rev. Lett.*, **124**, 170501, 2020.
- [99] K. C. MIAO, M. McEWEN, J. ATALAYA, D. KAFRI, L. P. PRYADKO, A. BENGTSSON, A. OPREMCÁK, K. J. SATZINGER, Z. CHEN, P. V. KLIMOV, *et al.* Overcoming leakage in quantum error correction. *Nature Physics*, pages 1–7, 2023.
- [100] M. GRASSL, T. BETH, and T. PELLIZZARI. *Codes for the quantum erasure channel*. *Phys. Rev. A*, **56**, 33–38, 1997.
- [101] C. H. BENNETT, D. P. DIVINCENZO, and J. A. SMOLIN. *Capacities of quantum erasure channels*. *Phys. Rev. Lett.*, **78**, 3217–3220, 1997.
- [102] T. M. STACE, S. D. BARRETT, and A. C. DOHERTY. *Thresholds for topological codes in the presence of loss*. *Phys. Rev. Lett.*, **102**, 200501, 2009.
- [103] S. D. BARRETT and T. M. STACE. *Fault tolerant quantum computation with very high threshold for loss errors*. *Phys. Rev. Lett.*, **105**, 200502, 2010.
- [104] A. KUBICA, A. HAIM, Y. VAKNIN, F. BRANDÃO, and A. RETZKER. *Erasure qubits: Overcoming the t_1 limit in superconducting circuits*, 2022.
- [105] Y. WU, S. KOLKOWITZ, S. PURI, and J. D. THOMPSON. Erasure conversion for fault-tolerant quantum computing in alkaline earth rydberg atom arrays. *Nature Communications*, **13** (1), 4657, 2022.
- [106] J. F. MARQUES, H. ALI, B. M. VARBANOV, M. FINKEL, H. M. VEEN, S. L. M. VAN DER MEER, S. VALLES-SANCLEMENTE, N. MUTHUSUBRAMANIAN, M. BEEKMAN, N. HAIDER, B. M. TERHAL, and L. DICARLO. *All-microwave leakage reduction units for quantum error correction with superconducting transmon qubits*. *Phys. Rev. Lett.*, **130**, 250602, 2023.
- [107] N. LACROIX, L. HOFELE, A. REMM, O. BENHAYOUNE-KHADRAOUI, A. McDONALD, R. SHILLITO, S. LAZAR, C. HELLINGS, F. SWIADEK, D. COLAO-ZANUZ, *et al.* Fast flux-activated leakage reduction for superconducting quantum circuits. *arXiv preprint arXiv:2309.07060*, 2023.
- [108] S. SPITZ, B. M. TARASINSKI, C. BEENAKKER, and T. O'BRIEN. *Adaptive weight estimator for quantum error correction in a time-dependent environment*. *Advanced Quantum Technologies*, **1** (1), 1800012, 2017.
- [109] E. H. CHEN, T. J. YODER, Y. KIM, N. SUNDARESAN, S. SRINIVASAN, M. LI, *et al.* Calibrated decoders for experimental quantum error correction. *Phys. Rev. Lett.*, **128** (11), 110504, 2022.
- [110] D. G. CORY, M. D. PRICE, W. MAAS, E. KNILL, R. LAFLAMME, W. H. ZUREK, T. F. HAVEL, and S. S. SOMAROO. Experimental quantum error correction. *Phys. Rev. Lett.*, **81**, 2152, 1998.
- [111] J. CHIAVERINI, D. LEIBFRIED, T. SCHAEZT, M. D. BARRETT, R. B. BLAKESTAD, J. BRITTON, W. M. ITANO, J. D. JOST, E. KNILL, C. LANGER, R. OZERI, and D. J. WINELAND. Realization of quantum error correction. *Nature*, **432** (7017), 602, 2004.
- [112] P. SCHINDLER, J. T. BARREIRO, T. MONZ, V. NEBENDAHL, D. NIGG, M. CHWALLA, M. HENNRICH, and R. BLATT. Experimental repetitive quantum error correction. *Science*, **332** (6033), 1059, 2011.

- [113] N. M. LINKE, M. GUTIERREZ, K. A. LANDSMAN, C. FIGGATT, S. DEBNATH, K. R. BROWN, and C. MONROE. Fault-tolerant quantum error detection. *Science advances*, **3** (10), e1701074, 2017.
- [114] C. RYAN-ANDERSON, N. BROWN, M. ALLMAN, B. ARKIN, G. ASA-ATTUAH, C. BALDWIN, J. BERG, J. BOHNET, S. BRAXTON, N. BURDICK, *et al.* Implementing fault-tolerant entangling gates on the five-qubit code and the color code. *arXiv preprint arXiv:2208.01863*, 2022.
- [115] T. B. PITTMAN, B. C. JACOBS, and J. D. FRANSON. Demonstration of quantum error correction using linear optics. *Phys. Rev. A*, **71**, 052332, 2005.
- [116] X.-C. YAO, T.-X. WANG, H.-Z. CHEN, W.-B. GAO, A. G. FOWLER, R. RAUSSENDORF, Z.-B. CHEN, N.-L. LIU, C.-Y. LU, Y.-J. DENG, *et al.* Experimental demonstration of topological error correction. *Nature*, **482** (7386), 489–494, 2012.
- [117] B. A. BELL, D. A. HERRERA-MARTÍ, M. S. TAME, D. MARKHAM, W. J. WADSWORTH, and J. G. RARITY. *Experimental demonstration of a graph state quantum error-correction code*. *Nat. Comm.*, **5**, 3658, 2014.
- [118] C. VIGLIAR, S. PAESANI, Y. DING, J. C. ADCOCK, J. WANG, S. MORLEY-SHORT, D. BACCO, L. K. OXENLÖWE, M. G. THOMPSON, J. G. RARITY, *et al.* Error-protected qubits in a silicon photonic chip. *Nature Physics*, **17** (10), 1137–1143, 2021.
- [119] G. WALDHERR, Y. WANG, S. ZAISER, M. JAMALI, T. SCHULTE-HERBRÜGGEN, H. ABE, T. OHSHIMA, J. ISOYA, J. F. DU, P. NEUMANN, and J. WRACHTRUP. Quantum error correction in a solid-state hybrid spin register. *Nature*, **506** (7487), 204, 2014.
- [120] T. H. TAMINIAU, J. CRAMER, T. VAN DER SAR, V. V. DOBROVITSKI, and R. HANSON. Universal control and error correction in multi-qubit spin registers in diamond. *Nat. Nanotech.*, **9** (3), 171, 2014.
- [121] M. ABOBEIH, Y. WANG, J. RANDALL, S. LOENEN, C. BRADLEY, M. MARKHAM, D. TWITCHEN, B. TERHAL, and T. TAMINIAU. Fault-tolerant operation of a logical qubit in a diamond quantum processor. *Nature*, **606** (7916), 884–889, 2022.
- [122] D. BLUVSTEIN, H. LEVINE, G. SEMEGHINI, T. T. WANG, S. EBADI, M. KALINOWSKI, A. KEESLING, N. MASKARA, H. PICHLER, M. GREINER, *et al.* A quantum processor based on coherent transport of entangled atom arrays. *Nature*, **604** (7906), 451–456, 2022.
- [123] D. BLUVSTEIN, S. J. EVERED, A. A. GEIM, S. H. LI, H. ZHOU, T. MANOVITZ, S. EBADI, M. CAIN, M. KALINOWSKI, D. HANGLEITER, *et al.* Logical quantum processor based on reconfigurable atom arrays. *Nature*, pages 1–3, 2023.
- [124] M. D. REED, L. DICARLO, S. E. NIGG, L. SUN, L. FRUNZIO, S. M. GIRVIN, and R. J. SCHOELKOPF. Realization of three-qubit quantum error correction with superconducting circuits. *Nature*, **482**, 382, 2012.
- [125] A. D. CÓRCOLES, E. MAGESAN, S. J. SRINIVASAN, A. W. CROSS, M. STEFFEN, J. M. GAMBETTA, and J. M. CHOW. *Demonstration of a quantum error detection code using a square lattice of four superconducting qubits*. *Nat. Comm.*, **6**, 6979, 2015.
- [126] J. F. MARQUES, B. M. VARBANOV, M. S. MOREIRA, H. ALI, N. MUTHUSUBRAMANIAN, C. ZACHARIADIS, F. BATTISTEL, M. BEEKMAN, N. HAIDER, W. VLOTHUIZEN, A. BRUNO, B. M. TERHAL, and L. DICARLO. *Logical-qubit operations in an error-detecting surface code*. *Nat. Phys.*, **18** (1), 80–86, 2022.
- [127] J. M. CHOW, J. M. GAMBETTA, E. MAGESAN, D. W. ABRAHAM, A. W. CROSS, B. R. JOHNSON, N. A. MASLUK, C. A. RYAN, J. A. SMOLIN, S. J. SRINIVASAN, and M. STEFFEN. *Implementing a strand of a scalable fault-tolerant quantum computing fabric*. *Nat. Comm.*, **5**, 4015, 2014.

- [128] E. JEFFREY, D. SANK, J. Y. MUTUS, T. C. WHITE, J. KELLY, R. BARENDs, Y. CHEN, Z. CHEN, B. CHIARO, A. DUNSWORTH, A. MEGRANT, P. J. J. O'MALLEY, C. NEILL, *et al.* *Fast accurate state measurement with superconducting qubits.* Phys. Rev. Lett., **112**, 190 504, 2014.
- [129] N. T. BRONN, E. MAGESAN, N. A. MASLUK, J. M. CHOW, J. M. GAMBETTA, and M. STEFFEN. *Reducing spontaneous emission in circuit quantum electrodynamics by a combined readout/filter technique.* IEEE T. Appl. Supercon., **25** (5), 1–10, 2015.
- [130] C. C. BULTINK, M. A. ROL, T. E. O'BRIEN, X. FU, B. C. S. DIKKEN, C. DICKEL, R. F. L. VERMEULEN, J. C. DE STERKE, A. BRUNO, R. N. SCHOUTEN, and L. DICARLO. *Active resonator reset in the nonlinear dispersive regime of circuit QED.* Phys. Rev. Appl., **6**, 034 008, 2016.
- [131] A. Y. KITAEV. Fault-tolerant quantum computation by anyons. Annals of Physics, **303** (1), 2–30, 2003.
- [132] E. DENNIS, A. KITAEV, A. LANDAHL, and J. PRESKILL. *Topological quantum memory.* Journal of Mathematical Physics, **43** (4452), 2002.
- [133] R. RAUSSENDORF and J. HARRINGTON. *Fault-tolerant quantum computation with high threshold in two dimensions.* Phys. Rev. Lett., **98**, 190 504, 2007.
- [134] T. THORBECK, A. EDDINS, I. LAUER, D. T. MCCLURE, and M. CARROLL. *Two-level-system dynamics in a superconducting qubit due to background ionizing radiation.* PRX Quantum, **4**, 020 356, 2023.
- [135] S. KRINNER, S. LAZAR, A. REMM, C. K. ANDERSEN, N. LACROIX, G. J. NORRIS, *et al.* *Benchmarking coherent errors in controlled-phase gates due to spectator qubits.* Phys. Rev. Appl., **14**, 024 042, 2020.
- [136] M. C. COLLODO, J. HERRMANN, N. LACROIX, C. K. ANDERSEN, A. REMM, S. LAZAR, J.-C. BESSE, T. WALTER, A. WALLRAFF, and C. EICHLER. *Implementation of conditional phase gates based on tunable zz interactions.* Phys. Rev. Lett., **125**, 240 502, 2020.
- [137] E. A. SETE, A. Q. CHEN, R. MANENTI, S. KULSHRESHTHA, and S. POLETO. *Floating tunable coupler for scalable quantum computing architectures.* Phys. Rev. Appl., **15**, 064 063, 2021.
- [138] F. MARXER, A. VEPSÄLÄINEN, S. W. JOLIN, J. TUORILA, A. LANDRA, C. OCKELOEN-KORPPI, W. LIU, O. AHONEN, A. AUER, L. BELZANE, V. BERGHOLM, C. F. CHAN, K. W. CHAN, *et al.* *Long-distance transmon coupler with cz-gate fidelity above 99.8%.* PRX Quantum, **4**, 010 314, 2023.
- [139] M. FIELD, A. Q. CHEN, B. SCHARMANN, E. A. SETE, F. ORUC, K. VU, V. KOSENKO, J. Y. MUTUS, S. POLETO, and A. BESTWICK. *Modular superconducting qubit architecture with a multi-chip tunable coupler.* arXiv preprint arXiv:2308.09240, 2023.
- [140] E. T. CAMPBELL, H. ANWAR, and D. E. BROWNE. *Magic-state distillation in all prime dimensions using quantum reed-muller codes.* Physical Review X, **2** (4), 041 021, 2012.
- [141] N. C. JONES, J. D. WHITFIELD, P. L. MCMAHON, M.-H. YUNG, R. VAN METER, A. ASPURUGUZI, and Y. YAMAMOTO. *Simulating chemistry efficiently on fault-tolerant quantum computers.* arXiv preprint arXiv:1204.0567, 2012.
- [142] A. G. FOWLER, S. J. DEVITT, and C. JONES. *Surface code implementation of block code state distillation.* Scientific reports, **3** (1), 1939, 2013.
- [143] C. GIDNEY and M. EKERÅ. *How to factor 2048 bit RSA integers in 8 hours using 20 million noisy qubits.* Quantum, **5**, 433, 2021.
- [144] C. NAYAK, S. H. SIMON, A. STERN, M. FREEDMAN, and S. DAS SARMA. *Non-abelian anyons and topological quantum computation.* Rev. Mod. Phys., **80**, 1083–1159, 2008.

- [145] M. LEIJNSE and K. FLENSBERG. Introduction to topological superconductivity and majorana fermions. *Semiconductor Science and Technology*, **27** (12), 124 003, 2012.
- [146] Y. KIM, A. EDDINS, S. ANAND, K. X. WEI, E. VAN DEN BERG, S. ROSENBLATT, H. NAYFEH, Y. WU, M. ZALETEL, K. TEMME, *et al.* Evidence for the utility of quantum computing before fault tolerance. *Nature*, **618** (7965), 500–505, 2023.
- [147] V. SIVAK, A. EICKBUSCH, B. ROYER, S. SINGH, I. TSIOUTSIOS, S. GANJAM, A. MIANO, B. BROCK, A. DING, L. FRUNZIO, *et al.* Real-time quantum error correction beyond break-even. *Nature*, **616** (7955), 50–55, 2023.
- [148] H. LEVINE, A. HAIM, J. S. HUNG, N. ALIDOUST, M. KALAEI, L. DELORENZO, E. A. WOLLACK, P. A. ARRIOLA, A. KHALAJHEDAYATI, Y. VAKNIN, *et al.* Demonstrating a long-coherence dual-rail erasure qubit using tunable transmons. *arXiv preprint arXiv:2307.08737*, 2023.
- [149] D. YOST, M. SCHWARTZ, J. MALLEK, D. ROSENBERG, C. STULL, J. YODER, G. CALUSINE, M. COOK, R. DAS, A. L. DAY, E. B. GOLDEN, D. K. KIM, A. MELVILLE, *et al.* *Solid-state qubits integrated with superconducting through-silicon vias*. *npj Quantum Inf.*, **6** (1), 1–7, 2020.
- [150] J. A. ALFARO-BARRANTES, M. MASTRANGELI, D. J. THOEN, S. VISSER, J. BUENO, J. J. A. BASELMANS, and P. M. SARRO. Superconducting high-aspect ratio through-silicon vias with DC-sputtered Al for quantum 3D integration. *IEEE Electron Device Lett.*, **41** (7), 1114–1117, 2020.
- [151] J. MALLEK, D. YOST, D. ROSENBERG, J. YODER, G. CALUSINE, M. COOK, *et al.* *Fabrication of superconducting through-silicon vias*. *arXiv:2103.08536*, 2021.
- [152] T. M. HAZARD, W. WOODS, D. ROSENBERG, R. DAS, C. F. HIRJIBEHEDIN, D. K. KIM, J. KNECHT, J. MALLEK, A. MELVILLE, B. M. NIEDZIELSKI, *et al.* Characterization of superconducting through-silicon vias as capacitive elements in quantum circuits. *Applied Physics Letters*, **123** (15), 2023.
- [153] G. J. NORRIS, L. MICHAUD, D. PAHL, M. KERSCHBAUM, C. EICHLER, J.-C. BESSE, and A. WALL-RAFF. Improved parameter targeting in 3d-integrated superconducting circuits through a polymer spacer process. *arXiv preprint arXiv:2307.00046*, 2023.
- [154] C. DICKEL, J. J. WESDORP, N. K. LANGFORD, S. PEITER, R. SAGASTIZABAL, A. BRUNO, B. CRIGER, F. MOTZOI, and L. DICARLO. *Chip-to-chip entanglement of transmon qubits using engineered measurement fields*. *Phys. Rev. B*, **97**, 064 508, 2018.
- [155] A. GOLD, J. PAQUETTE, A. STOCKKLAUSER, M. J. REAGOR, M. S. ALAM, A. BESTWICK, N. DIER, A. NERSISYAN, F. ORUC, A. RAZAVI, *et al.* Entanglement across separate silicon dies in a modular superconducting qubit device. *npj Quantum Information*, **7** (1), 142, 2021.
- [156] S. STORZ, J. SCHÄR, A. KULIKOV, P. MAGNARD, P. KURPIERS, J. LÜTOLF, T. WALTER, A. COPETUDO, K. REUER, A. AKIN, *et al.* Loophole-free bell inequality violation with superconducting circuits. *Nature*, **617** (7960), 265–270, 2023.
- [157] V. E. MANUCHARYAN, J. KOCH, L. I. GLAZMAN, and M. H. DEVORET. *Fluxonium: Single Cooper-pair circuit free of charge offsets*. *Science*, **326** (5949), 113–116, 2009. <http://www.sciencemag.org/content/326/5949/113.full.pdf>.
- [158] N. EARNEST, S. CHAKRAM, Y. LU, N. IRONS, R. K. NAIK, N. LEUNG, L. OCOLA, D. A. CZAPLEWSKI, B. BAKER, J. LAWRENCE, J. KOCH, and D. I. SCHUSTER. *Realization of a Λ system with metastable states of a capacitively shunted fluxonium*. *Phys. Rev. Lett.*, **120**, 150 504, 2018.
- [159] L. B. NGUYEN, Y.-H. LIN, A. SOMOROFF, R. MENCIA, N. GRABON, and V. E. MANUCHARYAN. *High-coherence fluxonium qubit*. *Phys. Rev. X*, **9**, 041 041, 2019.

- [160] L. B. NGUYEN, G. KOOLSTRA, Y. KIM, A. MORVAN, T. CHISTOLINI, S. SINGH, K. N. NESTEROV, C. JÜNGER, L. CHEN, Z. PEDRAMRAZI, B. K. MITCHELL, J. M. KREIKEBAUM, S. PURI, *et al.* *Blueprint for a high-performance fluxonium quantum processor*. PRX Quantum, **3**, 037 001, 2022.
- [161] A. SOMOROFF, Q. FICHEUX, R. A. MENCIA, H. XIONG, R. KUZMIN, and V. E. MANUCHARYAN. *Millisecond coherence in a superconducting qubit*. Phys. Rev. Lett., **130**, 267 001, 2023.
- [162] L. DING, M. HAYS, Y. SUNG, B. KANNAN, J. AN, A. DI PAOLO, A. H. KARAMLOU, T. M. HAZARD, K. AZAR, D. K. KIM, B. M. NIEDZIELSKI, A. MELVILLE, M. E. SCHWARTZ, *et al.* *High-fidelity, frequency-flexible two-qubit fluxonium gates with a transmon coupler*. Phys. Rev. X, **13**, 031 035, 2023.
- [163] N. P. BREUCKMANN and J. N. EBERHARDT. *Quantum low-density parity-check codes*. PRX Quantum, **2**, 040 101, 2021.
- [164] L. Z. COHEN, I. H. KIM, S. D. BARTLETT, and B. J. BROWN. *Low-overhead fault-tolerant quantum computing using long-range connectivity*. Science Advances, **8** (20), eabn1717, 2022.
- [165] S. BRAVYI, A. W. CROSS, J. M. GAMBETTA, D. MASLOV, P. RALL, and T. J. YODER. *High-threshold and low-overhead fault-tolerant quantum memory*. arXiv preprint arXiv:2308.07915, 2023.
- [166] E. T. CAMPBELL, B. M. TERHAL, and C. VUILLOT. *Roads towards fault-tolerant universal quantum computation*. Nature, **549** (7671), 172–179, 2017.
- [167] D. LITINSKI. *A Game of Surface Codes: Large-Scale Quantum Computing with Lattice Surgery*. Quantum, **3**, 128, 2019.
- [168] C. HORSMAN, A. G. FOWLER, S. DEVITT, and R. V. METER. *Surface code quantum computing by lattice surgery*. New J. Phys., **14** (12), 123 011, 2012.

