

Demonstrating a Long-Coherence Dual-Rail Erasure Qubit Using Tunable Transmons

H. Levine¹, A. Haim,^{1,2} J. S. C. Hung,¹ N. Alidoust,¹ M. Kalae,¹ L. DeLorenzo,¹ E. A. Wollack¹,
 P. Arrangoiz-Arriola,¹ A. Khalajhedayati¹, R. Sanil,¹ H. Moradinejad,¹ Y. Vaknin,^{1,3} A. Kubica¹, D. Hover,¹
 S. Aghaeimeibodi¹, J. A. Alcid,¹ C. Baek,¹ J. Barnett,¹ K. Bawdekar,¹ P. Bienias,¹ H. A. Carson,¹ C. Chen,¹
 L. Chen¹, H. Chinkezan,¹ E. M. Chisholm,¹ A. Clifford,¹ R. Cosmic,¹ N. Crisosto,¹ A. M. Dalzell,¹ E. Davis,¹
 J. M. D'Ewart,¹ S. Diez,¹ N. D'Souza,¹ P. T. Dumitrescu¹, E. Elkhoully,¹ M. T. Fang,¹ Y. Fang¹, S. Flammia,¹
 M. J. Fling,¹ G. Garcia,¹ M. K. Gharzai,¹ A. V. Gorshkov¹, M. J. Gray,¹ S. Grimberg,¹ A. L. Grimsom^{1,4},
 C. T. Hann,¹ Y. He,¹ S. Heidel^{1,†}, S. Howell,¹ M. Hunt,¹ J. Iverson,¹ I. Jarrige,¹ L. Jiang,^{1,5} W. M. Jones¹,
 R. Karabalin,¹ P. J. Karalekas¹, A. J. Keller¹, D. Lasi,¹ M. Lee¹, V. Ly,¹ G. MacCabe,¹ N. Mahuli,¹ G. Marcaud¹,
 M. H. Matheny,¹ S. McArdle,¹ G. McCabe¹, G. Merton,¹ C. Miles,¹ A. Milsted,¹ A. Mishra¹, L. Monceli¹,
 M. Naghiloo,¹ K. Noh,¹ E. Oblepias,¹ G. Ortuno,¹ J. C. Owens,¹ J. Pagdilao,¹ A. Panduro,¹ J.-P. Paquette,¹ R. N. Patel,¹
 G. Peairs,¹ D. J. Perello,¹ E. C. Peterson,¹ S. Ponte,¹ H. Putterman,¹ G. Refael,^{1,6,7} P. Reinhold,¹ R. Resnick,¹
 O. A. Reyna,¹ R. Rodriguez,¹ J. Rose,¹ A. H. Rubin^{1,‡}, M. Runyan,¹ C. A. Ryan¹, A. Sahmoud,¹ T. Scaffidi,^{1,§}
 B. Shah,¹ S. Siavoshi,¹ P. Sivarajah,¹ T. Skogland,¹ C.-J. Su,¹ L. J. Swenson,¹ J. Sylvia,¹ S. M. Teo,¹
 A. Tomada,¹ G. Torlai,¹ M. Wistrom,^{1,||} K. Zhang,¹ I. Zuk^{1,3}, A. A. Clerk,^{1,5} F. G. S. L. Brandão,^{1,6}
 A. Retzker,^{1,3} and O. Painter^{1,6,8,*}

¹AWS Center for Quantum Computing, Pasadena, California 91125, USA

²Institute of Applied Physics, The Hebrew University of Jerusalem, Jerusalem 91904, Givat Ram, Israel

³Racah Institute of Physics, The Hebrew University of Jerusalem, Jerusalem 91904, Givat Ram, Israel

⁴Centre for Engineered Quantum Systems, School of Physics, The University of Sydney, Sydney, New South Wales 2006, Australia

⁵Pritzker School of Molecular Engineering, University of Chicago, Chicago, Illinois 60637, USA

⁶Institute for Quantum Information and Matter, California Institute of Technology, Pasadena, California 91125, USA

⁷Department of Physics, California Institute of Technology, Pasadena, California 91125, USA

⁸Thomas J. Watson, Sr., Laboratory of Applied Physics and Kavli Nanoscience Institute, California Institute of Technology, Pasadena, California 91125, USA



(Received 28 September 2023; revised 27 December 2023; accepted 16 January 2024; published 20 March 2024)

Quantum error correction with erasure qubits promises significant advantages over standard error correction due to favorable thresholds for erasure errors. To realize this advantage in practice requires a qubit for which nearly all errors are such erasure errors, and the ability to check for erasure errors without dephasing the qubit. We demonstrate that a “dual-rail qubit” consisting of a pair of resonantly coupled transmons can form a highly coherent erasure qubit, where transmon T_1 errors are converted into erasure errors and residual dephasing is strongly suppressed, leading to millisecond-scale coherence within the qubit subspace. We show that single-qubit gates are limited primarily by erasure errors, with erasure probability $p_{\text{erasure}} = 2.19(2) \times 10^{-3}$ per gate while the residual errors are ~ 40 times lower. We further demonstrate midcircuit detection of erasure errors while introducing $< 0.1\%$ dephasing error per check. Finally, we show that the suppression of transmon noise allows this dual-rail qubit to preserve high coherence over a broad tunable operating range, offering an improved capacity to avoid frequency collisions. This work establishes transmon-based dual-rail qubits as an attractive building block for hardware-efficient quantum error correction.

DOI: [10.1103/PhysRevX.14.011051](https://doi.org/10.1103/PhysRevX.14.011051)

Subject Areas: Quantum Physics, Quantum Information

*Corresponding author: ojp@amazon.com

[†]Present address: OpenAI, San Francisco, California, USA.

[‡]Present address: Department of Physics and Electrical and Computer Engineering, University of California, Davis, California 95616, USA.

[§]Present address: Department of Physics and Astronomy, University of California, Irvine, California 92697, USA.

^{||}Present address: Google, 1600 Amphitheatre Parkway, Mountain View, California 94043, USA.

Published by the American Physical Society under the terms of the [Creative Commons Attribution 4.0 International](https://creativecommons.org/licenses/by/4.0/) license. Further distribution of this work must maintain attribution to the author(s) and the published article's title, journal citation, and DOI.

I. INTRODUCTION

The path toward useful quantum error correction requires qubits which have physical error rates well below the error correction threshold. While many physical qubit platforms have reached error rates below the surface code threshold [1–5], a major frontier challenge is to continue driving down error rates to reduce the overhead necessary to protect logical qubits. In parallel with this effort, the growing field of hardware-efficient error correction seeks to leverage or engineer the noise properties of physical qubits to more efficiently correct errors, effectively raising thresholds and reducing overhead [6–10]. Recent examples include noise-bias engineering and robust qubit encoding in superconducting bosonic modes [11,12].

The recently developed *erasure qubit* offers a compelling path toward more relaxed error correction requirements [13–16]. While most physical qubits exhibit undetectable errors that occur within the qubit subspace, erasure qubits are those in which the dominant error is leakage out of the computational subspace, and for which these leakage errors can be detected in real time. These *erasure errors* can be corrected more efficiently, leading to higher thresholds and reduced logical error rates [13–15,17–19]. To leverage this advantage, erasure qubits must demonstrate (1) a large erasure noise bias, which is the ratio of erasure errors to residual errors within the qubit subspace, and (2) the ability to detect erasure errors without introducing additional errors within the subspace. While erasure qubit implementations have been proposed in several architectures including neutral atoms [13,19], ions [20,21], and superconducting qubits [14,16], these ideas have only recently been implemented experimentally using neutral atoms, demonstrating that a large fraction of gate errors can be detected midcircuit and converted into erasures [22], as well as the excision of state preparation and Rydberg decay errors detected as erasures [23].

In this paper, we demonstrate that a superconducting “dual-rail qubit” formed by two resonantly coupled transmons meets both requirements for serving as an erasure qubit. We show that the dual-rail qubit converts transmon T_1 errors into detectable erasure errors while passively suppressing residual transmon dephasing, enabling a large erasure noise bias while idling and during single-qubit gates. We then demonstrate high-fidelity midcircuit erasure detection using an ancilla qubit, and show that this detection minimally degrades coherence within the dual-rail subspace. Finally, we demonstrate that the dual-rail qubit’s passive noise suppression allows it to maintain high coherence while remaining broadly tunable, offering an extra knob to dodge frequency collisions and improving the prospects for qubit yield at scale.

II. DUAL-RAIL QUBIT

The dual-rail qubit is defined in the single-excitation manifold of a pair of superconducting modes [14,16,24–28].

We use a pair of resonantly coupled transmons, as first demonstrated in Ref. [28], where the logical subspace consists of the hybridized symmetric and antisymmetric states $|0_L\rangle = (1/\sqrt{2})(|01\rangle - |10\rangle)$ and $|1_L\rangle = (1/\sqrt{2})(|01\rangle + |10\rangle)$. This encoding allows T_1 decay of the underlying transmons, which constitutes the fundamental limit to transmon coherence, to be converted to erasure errors in the form of leakage to the $|00\rangle$ state [Figs. 1(a)–1(c)].

Crucially, in addition to converting T_1 errors into erasure errors, the resonant coupling between $|01\rangle$ and $|10\rangle$ with strength g acts as a passive decoupling mechanism which strongly suppresses the impact of low-frequency noise on the underlying transmons [14,28], analogous to continuous dynamical decoupling in driven systems [29–31]. This is illustrated by the dual-rail energy gap $E_{\text{DR}} \approx \sqrt{(2g)^2 + \delta^2}$, where δ is the frequency difference between the two transmons which inherits frequency noise from each transmon. When $g \gg \delta$, frequency noise is suppressed according to $E_{\text{DR}} \approx 2g + \delta^2/4g$, allowing the dual-rail energy gap to be orders of magnitude more stable than the underlying transmons (see Appendix M for a more detailed analysis). This noise suppression enables a large erasure noise bias, in which dephasing errors within the subspace are suppressed and the dominant error mechanism is erasure errors due to the underlying transmon T_1 .

Our experiments utilize a superconducting quantum circuit with three tunable transmons, two of which (Q1 and Q2) form the dual-rail qubit and are parked on resonance with one another, while the third (Q3) is used as an ancilla qubit (Fig. 1). Experimental sequences involving dual-rail erasure qubits are illustrated in Fig. 1(d), with calibration routines presented in Appendices D and E. Each sequence begins with a microwave pulse on Q1 to initialize the logical $|1_L\rangle$. Single-qubit gates within the dual-rail subspace are performed by flux modulation of Q2 at the frequency $2g = 2\pi \times 180$ MHz, with the phase of the flux modulation defining the axis of the drive field. At the end of the circuit, the dual-rail qubit is measured by adiabatically separating the two transmons with a flux pulse on Q2 such that the two logical states $|0_L\rangle, |1_L\rangle$ map to the pair states $|01\rangle$ and $|10\rangle$, respectively, whereupon the two transmons are jointly read out [28]. Dual-rail qubit operation is illustrated in Fig. 1(f) by a spin-echo experiment, where the final readout shows both coherent oscillations within the $|0_L\rangle, |1_L\rangle$ subspace as well as leakage to $|00\rangle$ due to transmon T_1 decay.

During the experimental sequence, periodic erasure checks may be performed to identify if the system decayed into $|00\rangle$. This is done by leveraging a dispersive shift on the ancilla qubit which depends on whether the dual rail is in $|00\rangle$ or if it remains in the logical subspace [Fig. 1(e)]. The erasure check consists of a conditional π pulse on the ancilla which is resonant only if the dual rail is in $|00\rangle$, followed by ancilla state readout. We note that a similar procedure using a directly coupled readout resonator could

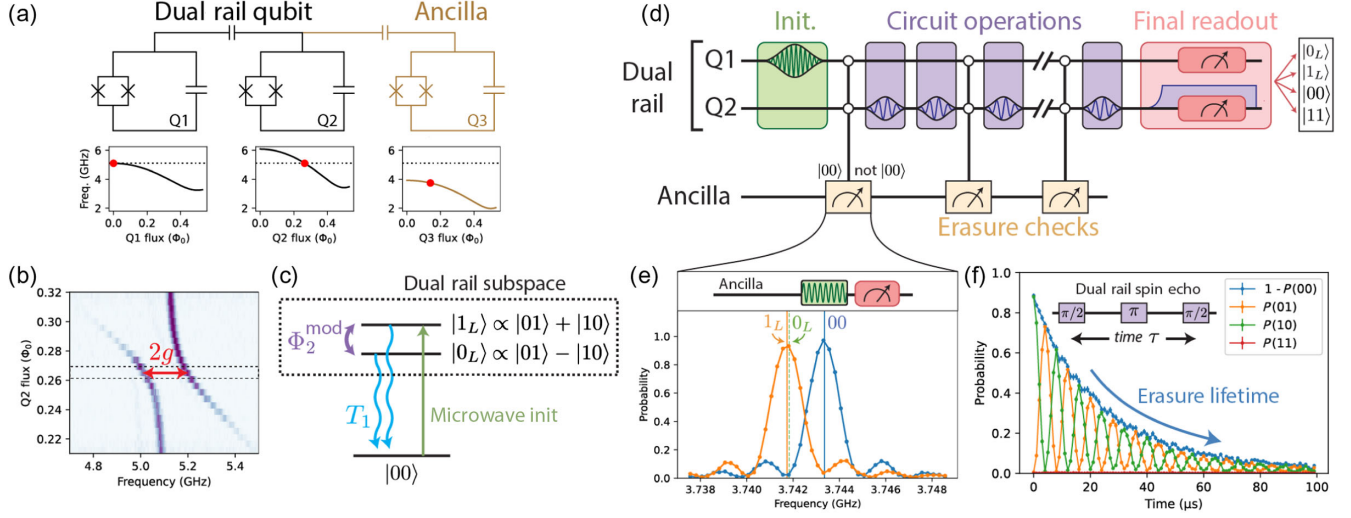


FIG. 1. Dual-rail erasure qubit encoding. (a) Our circuit has three tunable transmons: Q1 and Q2 form the dual-rail qubit, and Q3 is an ancilla qubit used for erasure detection (insets: tuning ranges, with primary operating point marked). (b) Microwave spectroscopy of Q1 as Q2 is tuned into resonance, exhibiting an avoided crossing with gap $2g/2\pi = 180$ MHz. (c) The dual-rail subspace is the symmetric and antisymmetric combination of $|01\rangle$ and $|10\rangle$. Decay of the underlying transmons leads to detectable leakage to $|00\rangle$, constituting an erasure error. (d) Circuits begin with a 40 ns microwave pulse to initialize $|1_L\rangle$, followed by single-qubit gates enacted by flux modulation of Q2, and finally the qubits are adiabatically separated and jointly read out. (e) We perform spectroscopy on the ancilla after initializing the dual-rail pair in $|00\rangle$ (blue), $|1_L\rangle$ (orange), or $|0_L\rangle$ (green, trace not shown). Erasure checks are performed by applying a microwave pulse on the ancilla which is resonant only if the dual-rail pair is in $|00\rangle$ (blue vertical line). (f) Spin-echo experiment on the dual-rail qubit where the phase of the final $\pi/2$ is stepped, showing both coherent fringes within the subspace and also decay out of the subspace with a lifetime of ~ 30 μ s. Error bars denote 68% confidence intervals unless otherwise specified.

also be used for erasure detection, but would require more carefully engineered symmetric coupling strengths to avoid dephasing as discussed in Ref. [14]. Such erasure checks are critical tools when using erasure qubits for quantum error correction, and as discussed below, are also necessary for characterizing the coherence properties of the dual-rail qubit.

III. COHERENCE WITHIN THE DUAL-RAIL SUBSPACE

A key metric for erasure qubits is the erasure noise bias, which should be $\gg 1$ in order to benefit most significantly from erasure conversion [13]. This should hold for errors while idling, requiring that the coherence within the logical subspace is much longer than the erasure lifetime, as well as during gates. To probe error rates within the subspace, we perform coherence and gate benchmarking experiments and postselect on shots where the system stays in the dual-rail subspace. We note that postselection is used here only to characterize the nonerasure error rates and would not be needed in a concatenated surface code architecture as described in Ref. [14].

We postselect against leakage by relying on both the final readout of the dual-rail pair as well as a sequence of erasure checks performed during the measurement [Fig. 2(a)]. These midcircuit erasure checks are important to catch events in which the system decays into $|00\rangle$ during

the circuit but then heats back into the dual-rail subspace; without such checks, this decay and reheating process limits the dephasing time that can be measured using only the final readout. Erasure checks are evenly spaced during the coherence measurement and successfully mitigate the contributions of decay and reheating events as long as they are repeated with a spacing shorter than the transmon T_1 times (see further discussion in Appendix J).

We find that while the two transmons Q1 and Q2 individually have T_2^{CPMG} at the microsecond and tens of microseconds scales, the T_2^{CPMG} within the dual-rail subspace (postselected against erasure errors) is > 500 μ s, ranging from 543(23) μ s for one echo pulse to 1.25(8) ms for 256 echo pulses [Figs. 2(b) and 2(c)]. The dual-rail T_1 fits to an extrapolated $T_1 = 906(15)$ μ s [Fig. 2(d)]. These millisecond-scale coherence times are far longer than the timescale for decay out of the subspace into $|00\rangle$, which occurs at a characteristic erasure lifetime $T_{\text{erasure}} \sim 30$ μ s [Fig. 2(b)] set by the underlying transmon T_1 values 15–35 μ s (Table I). The ratio of these timescales defines the erasure noise bias for idling errors, $T_2^{\text{CPMG}}/T_{\text{erasure}} \gtrsim 20$, confirming that the dominant error on the dual-rail qubit while idling is erasure errors.

Several effects can limit the coherence within the dual-rail subspace at the millisecond level and are discussed in Appendix N. The most dominant contribution is thermal Johnson-Nyquist noise on the flux lines, which is expected

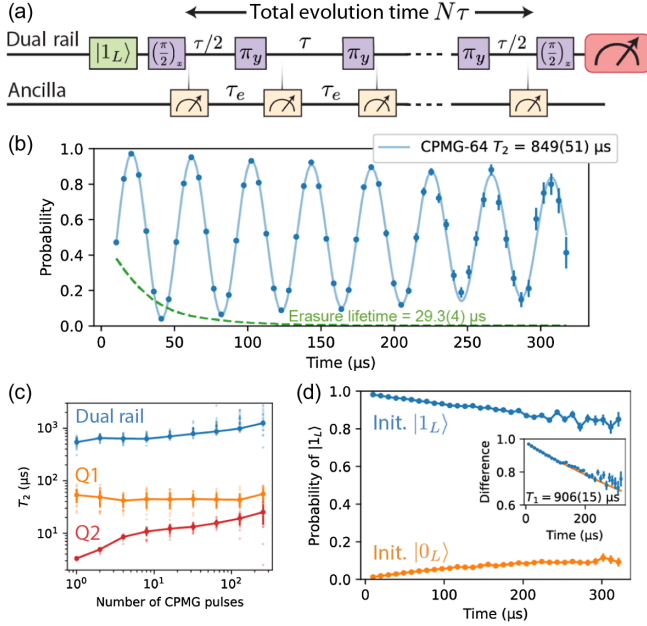


FIG. 2. Millisecond-scale coherence within the dual-rail subspace. (a) We measure T_2 coherence with a CPMG (Carr-Purcell-Meiboom-Gill) sequence [32] consisting of N π pulses with a total evolution time of $N\tau$. In parallel, we perform a sequence of $N_e = 11$ erasure checks, evenly spaced by $\tau_e = N\tau/(N_e - 1)$. Coherence data are postselected against leakage using both the final dual-rail readout and also the midcircuit erasure checks. (b) Example experiment with $N = 64$ π pulses, where the final $\pi/2$ pulse phase is stepped to produce fringes. The postselection probability decays with time according to a 29.3(4) μs erasure lifetime, while the postselected shots are fit to an exponential decay with extrapolated $T_2 = 849(51)$ μs . 330 000 shots were taken per point, but only a small fraction of these survive postselection. (c) CPMG measurements on each underlying transmon, if parked alone at the operating point, as well as on the dual-rail qubit. Error bars are standard deviation of many individual measurements. (d) T_1 within the dual-rail subspace. We measure the probability of ending in $|1_L\rangle$, postselected against erasure errors, after initializing in either $|1_L\rangle$ (blue) or $|0_L\rangle$ (orange). The difference between these traces is fit to an exponential decay with fixed offset 0, giving $T_1 = 906(15)$ μs . The 24 hour time trace of data comprising these dual-rail T_1 , T_2 plots is shown in Appendix L.

to limit the T_1 within the dual-rail subspace to ~ 1 ms. We therefore hypothesize that this is the limiting contribution for our measured T_1 ; this effect can be mitigated with further cryogenic attenuation of the fast-flux line and will be studied in future work. Other effects, including residual sensitivity to transmon frequency noise, photon fluctuations in readout resonators, and leakage into the two-photon manifold, can all limit dual-rail coherence at the level of several milliseconds. While the measured T_1 is attributed to Johnson noise, the T_2 may be limited by a combination of these effects and requires further investigation.

We find that the dual-rail T_2^* phase coherence without dynamical decoupling is limited by a separate effect: slow,

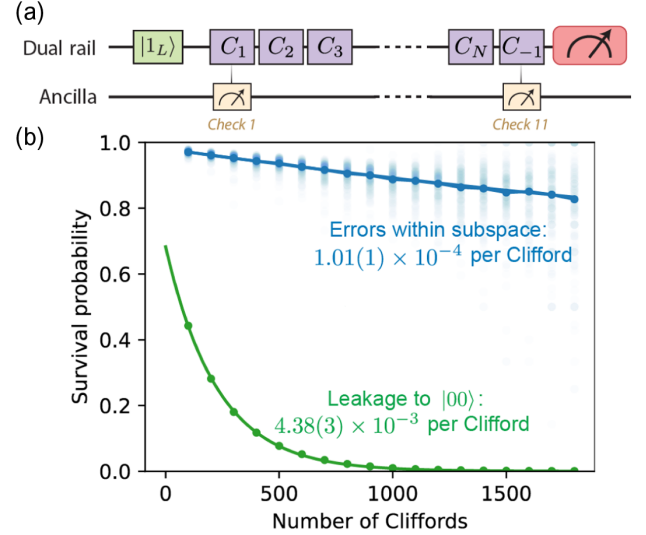


FIG. 3. Single-qubit randomized benchmarking. (a) The randomized benchmarking circuit uses random Cliffords, each implemented as two X90 pulses with appropriate phase shifts [33]. In parallel, we perform 11 erasure checks evenly spaced throughout the circuit, regardless of depth. The X90 pulse is a Gaussian-shaped 48 ns flux pulse on Q2 along with a 0.067 rad phase correction. An example timing diagram is presented in Fig. 9. (b) We plot the postselection probability (green) as a function of circuit depth and fit to an exponential decay to 0 to extract the erasure error per Clifford of $4.38(3) \times 10^{-3}$. The postselected survival probability within the dual-rail subspace (blue), averaged over 200 random circuits for each depth, is fit to an exponential decay with offset 1/2 and gives a residual error rate of $1.01(1) \times 10^{-4}$ per Clifford. Individual outcomes for random circuits are plotted with transparency. We symmetrize readout errors by alternating measurements in which the ideal outcome is $|0_L\rangle$ or $|1_L\rangle$. The X90 gate errors are half those of the Clifford gate, with erasure error $2.19(2) \times 10^{-3}$ and residual error $5.06(6) \times 10^{-5}$. Similar error rates are computed when postselecting only on midcircuit erasure checks (Appendix F).

seconds-scale telegraph noise on the dual-rail frequency (Fig. 13). We correlate this frequency instability with similar telegraph noise on one of the underlying transmons (Q1) and attribute this to a nearby toggling two-level system (TLS) defect at 4.98 GHz (Appendix K). This behavior is consistent with a dispersive coupling model discussed in Appendix N3 and may offer a new probe for TLS behavior. Nonetheless, the switching time is sufficiently slow that this noise is effectively mitigated with dynamical decoupling.

To validate that long dual-rail coherence times with dynamical decoupling translate to high-fidelity gates, we characterize single-qubit X90 gates on the dual-rail qubit, enacted by 48 ns flux-modulation pulses on Q2 (Appendix E). We perform Clifford randomized benchmarking (Fig. 3), in parallel with frequent erasure checks, and measure an erasure error probability of $2.19(2) \times 10^{-3}$ per X90, consistent with the duration of the gate relative to

the erasure lifetime. Postselected against erasures, we find a low residual error rate of $5.06(6) \times 10^{-5}$ per X90, averaged over a 24 hour period, with the remaining error likely limited by a combination of residual decoherence within the logical subspace and coherent calibration errors. These error rates constitute an erasure noise bias of 43(1), and show that single-qubit gates preserve the large noise bias achieved while idling.

IV. MIDCIRCUIT ERASURE DETECTION USING AN ANCILLA QUBIT

We next turn to the characterization of midcircuit erasure detection. Our erasure check relies on a dispersive shift ($2\pi \times 1.6$ MHz) induced on the ancilla qubit when the dual-rail qubit is in the logical subspace ($|0_L\rangle, |1_L\rangle$) relative to when it is in $|00\rangle$ [Fig. 1(e); precise shifts given in Appendix G] [34]. The erasure check consists of a 540 ns square microwave pulse on the ancilla qubit whose length is chosen to excite the ancilla only if the dual rail is in $|00\rangle$ and not if the dual rail is in its logical subspace [Fig. 1(e)]. After the microwave pulse, we perform a 340 ns readout of the ancilla transmon (Appendix B).

There are three metrics to evaluate the performance of erasure detection, motivated by the proposed use of erasure checks in an error-correcting code [13,14]. First, *false-positive* errors are those in which no error occurred but an erasure is falsely flagged. Second, *false-negative* errors are those in which an erasure error occurs but is not correctly flagged. Finally, *erasure-check-induced dephasing* is the dephasing on the dual-rail qubit induced by checking for an erasure error when no such error had occurred. To establish target performance levels for these metrics, we consider a surface code protocol in which an erasure check is performed alongside each 2-qubit gate to prevent the diffusion of errors. Erasure check errors will add to the effective error rates of the gates, and in instances where the erasure check is slower than the gates, an additional error will also be introduced due to the time associated with the erasure check. For concreteness, we benchmark against 1% erasure error per gate and 0.1% Pauli error per gate; this pair of error rates is comfortably below the surface code threshold [14].

False positives are characterized by initializing the dual rail in $|1_L\rangle$, performing an erasure check, and then reading out the dual-rail pair [Fig. 4(a)]. We postselect on the final pairwise readout showing the correct state was initialized, and measure the probability that the erasure check correctly indicated no error. We find a false-positive rate of 0.58(1)% for $|1_L\rangle$, limited by accidental excitation of the ancilla when the dual rail has not decayed. We find a higher false-positive rate of $\sim 0.8\%$ for the logical state $|+_L\rangle$ [measured independently in Fig. 4(c)], which we take as an effective average rate that does not distinguish between assignment errors and any mechanism which may induce an erasure error during the check. In a surface

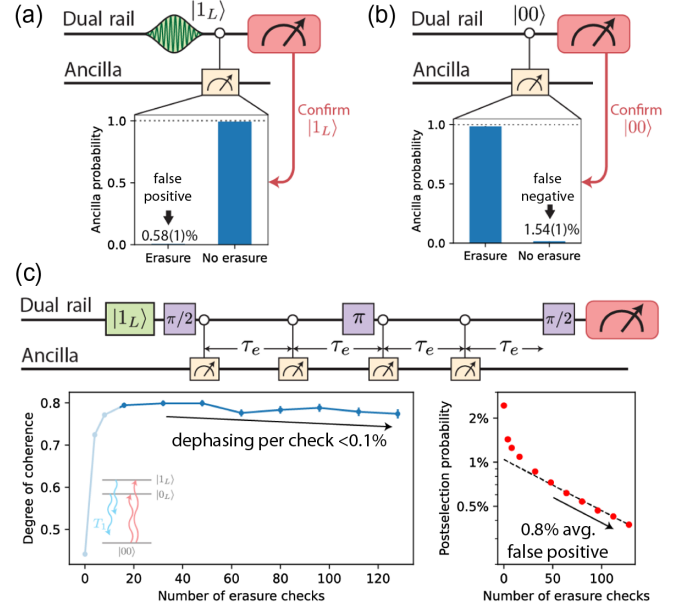


FIG. 4. Characterizing midcircuit erasure detection. (a),(b) We characterize false-positive and false-negative errors by initializing a target state, either $|1_L\rangle$ or $|00\rangle$, performing an erasure check, and then analyzing its results postselected on a final readout confirming the correct initialization. We ensure initialization of the ancilla in $|0\rangle$ by postselecting on an additional prereadout before the sequence begins (not shown). (c) We characterize erasure-check-induced dephasing by inserting a variable number of checks into a spin-echo experiment with fixed total evolution time $T = 114 \mu\text{s}$. The N erasure checks are evenly spaced with separation time $\tau_e = T/N$. We vary the phase of the final $\pi/2$ pulse to measure the remaining degree of coherence, postselected on all erasure checks indicating no erasure. Coherence appears to grow with a small number of checks due to properly catching and postselecting against leakage into $|00\rangle$ which reheats into the subspace. This effect saturates, after which we observe little dephasing per check at the level of $< 0.1\%$ (see Appendix I). The postselection probability (right-hand subplot) decreases with increasing erasure checks, first due to discarding shots with leakage, then saturating at a rate of $\sim 0.8\%$ per check which we interpret as a combination of false-positive assignment errors and check-induced erasure events. We note that a positive erasure detection event can with low probability, $0.0293(3)\%$, reexcite the dual-rail qubit (see Appendix B).

code context, this type of error would “inject” extra erasure errors into the system, and should thus be compared to the 1% erasure error target.

False negatives are characterized by initializing the dual rail in $|00\rangle$, performing an erasure check, and then postselecting on a final readout which shows $|00\rangle$ [Fig. 4(b)]. We find a false-negative rate of 1.54(1)%, limited by T_1 decay of the ancilla during readout. In a surface code, the quantity of interest is the probability that an erasure event occurred and was not properly flagged in detection, thus persisting and possibly propagating to undetectable Pauli errors. The probability of such an event is the chance of an erasure error during the erasure check,

$T_{\text{check}}/T_{\text{erasure}} \sim 2.9\%$, multiplied by the false negative rate. The probability of a missed erasure is then 0.04% , comfortably below the 0.1% target Pauli level.

Finally, we measure the dephasing induced by each erasure check by performing a spin-echo measurement on the dual rail and inserting a variable number of erasure checks [Fig. 4(c)]. The dual-rail coherence appears to improve and then plateau when inserting a small number of erasure checks, as these checks correctly eliminate shots in which the system decayed to $|00\rangle$ and then heated back into the subspace. As more checks are inserted, remaining phase coherence degrades only minimally, from which we compute a conservative upper bound of $< 0.1\%$ error per erasure check (Appendix I). Such dephasing errors inject undetectable Pauli errors in each check, and thus should be compared to the target 0.1% Pauli error level.

While the fidelities for these three metrics are below their respective surface code thresholds, the current erasure check time (880 ns) is slower than 2-qubit gate protocols in this architecture of ~ 200 ns [14,28], and would contribute a larger erasure error of $T_{\text{check}}/T_{\text{erasure}} \sim 2.9\%$ per gate. Faster erasure checks can be achieved through larger dispersive couplings and optimized ancilla readout [35], or by using a single symmetrically coupled readout resonator [14]. Additionally, ancilla reset and reinitialization of the dual rail after an erasure event are still needed, and we expect implementations can be closely adapted from standard transmon reset and leakage reduction protocols [36–38].

V. ROBUST OPERATION AT FLEXIBLE OPERATING POINTS

The dual-rail qubit enables a large erasure noise bias as well as high-fidelity erasure detection. Importantly, these features are not fine-tuned based on special operating frequencies, but instead can be preserved at a broad range of operating points. This robustness offers a significant advantage over standard architectures in which tunable transmons are expected to operate close to their flux-insensitive sweet spots, but may be unable to due to frequency collisions, parasitic modes, or TLSs. While operating tunable transmons away from the sweet spot is possible, extensive dynamical decoupling is needed to recover T_2 performance, and this decoupling may complicate single-qubit or multiqubit gates [39]. In the dual-rail qubit, however, this noise suppression is achieved passively, and carries over into the operation of single-qubit gates, erasure detection, and proposed implementations of multiqubit gates [14].

We test this robustness by parking the two dual-rail transmons at several operating points over a 350 MHz band from 4.75 to 5.1 GHz which is mutually accessible by each. We park each transmon individually at each of these operating points to first characterize their individual properties, and find that Q2 has fairly uniform coherence over this

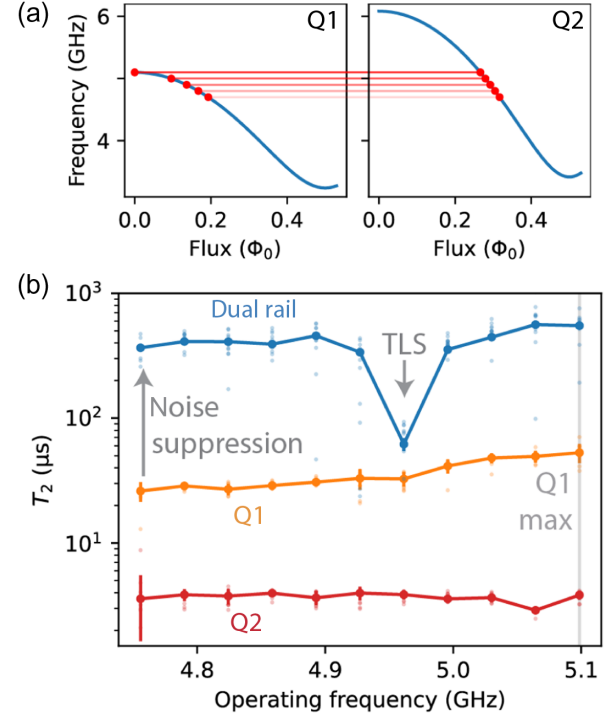


FIG. 5. Robust operation away from sweet spot. (a) The dual-rail qubit can be operated anywhere the two transmons can be brought onto resonance, forming an engineered “tunable” sweet spot. (b) At a range of operating points, we measure T_2^{echo} of each underlying transmon individually as well as of the dual-rail qubit. With the exception of a narrow region around 4.96 GHz, the dual-rail qubit maintains a coherence of several hundred microseconds over most of this range [$366(20)$ – $550(30)$ μs], with the modest degradation likely due to increased effects of Q1’s flux noise away from its sweet spot (Appendix N). These coherence times are over an order of magnitude larger than the higher-coherence transmon (Q1) and more than 2 orders of magnitude larger than Q2. The dip in dual-rail T_2 at 4.96 GHz is likely explained by a TLS coupled to Q2 which is near resonant with the upper hybrid mode at this operating point; a similar reduction is present on Q2 directly at an offset operating point, as visible in supplementary datasets presented in Appendix O. Error bars are standard deviations of many fit results (individual dots) for Q1 and Q2, and are fit uncertainties for the dual-rail qubit.

range since it is already far from its sweet spot, while Q1’s coherence degrades as it is tuned away from the sweet spot (Fig. 5). Across this full range, except for a narrow band at 4.96 GHz impacted by a TLS (discussed in Appendix O), the dual-rail qubit preserves T_2^{echo} coherence of hundreds of microseconds, from $366(20)$ to $550(30)$ μs. While more investigations are needed to fully understand coherence limitations, these measurements highlight the resilience of the dual-rail qubit to flux noise which normally precludes operating at highly flux-sensitive points. Additional adjustments of the qubit parameters, such as tuning ranges and coupling strengths, may improve performance and robustness further.

VI. DISCUSSION AND OUTLOOK

These results show that the dual-rail qubit is a promising, highly coherent building block for erasure-based error correction architectures. Coherence within the dual-rail subspace, even for transmons which are individually noisy, is approaching the millisecond regime, while the erasure lifetime is set by the transmon T_1 . Subsequent studies may complete the toolbox for operating an erasure surface code, including developing 2-qubit gates between dual-rail pairs [28], reinitializing dual-rail qubits after erasure errors, and exploring novel erasure detection mechanisms [14].

While the dual-rail transmon architecture is more complex than standard transmon architectures, an optimized approach can be taken which uses a single resonator for both readout and erasure detection [14]. In this setup, a dual-rail qubit involves only slightly more space on the processor and only one additional control line compared to standard transmons (Appendix P). For this modest increase in complexity, the conversion of T_1 decay to erasure errors enables leveraging of higher thresholds and effectively larger code distances as compared with Pauli errors, offering a positive prospect for accelerating the path to near-term below-threshold operation of error-corrected processors and for long-term logical qubit performance. Additionally, the broad effort toward scaling transmon systems and improving transmon coherence will translate naturally into the dual-rail architecture; for example, improvement in transmon T_1 will result in lower erasure error rates.

The passive noise suppression also opens opportunities to build dual-rail qubits out of other components which have long T_1 but shorter T_2 coherence times. Candidate components would include transmonlike circuits which operate at lower values of E_J/E_C , phase qubits, or potentially fluxonium. Such physical qubits, which in some regimes are undesirable when used on their own due to short coherence, may form highly robust logical qubits in a dual-rail architecture with improved erasure lifetime while still maintaining strong phase coherence within the dual-rail subspace.

Finally, in addition to prospects for enhancing logical qubit performance, we note that the ability of erasure qubits to postselect against errors may already offer significant advantages for some tasks including quantum simulation [23,40], sampling [41], and nonverifiable algorithms [40,42,43]. For these applications, postselecting on the lack of erasure errors enables more accurate reconstruction of the probability distributions than would be achieved by conventional quantum processors even with lower overall error rates [23].

Note added.—Recently, we became aware of related work on the dual-rail encoding in microwave cavities [44].

ACKNOWLEDGMENTS

We thank the technical support from across the AWS Center for Quantum Computing, including the teams involved with theory, design, fabrication, device packaging, cryogenics, signals, software, procurement, and lab infrastructure. We also thank Simone Severini, Bill Vass, and AWS for supporting the quantum computing program. Finally, we thank Jeff Thompson, Manuel Endres, and David Schuster for helpful discussions and feedback on the manuscript.

APPENDIX A: SUPERCONDUCTING DEVICE

The superconducting quantum circuit is shown in Fig. 6, and device parameters are summarized in Table I. Each of the three transmons has a dedicated readout resonator, all of which are coupled to a single transmission line. Q1 and Q3 have additional Purcell filters on their resonators. Transmon capacitive couplings are realized via wedge couplers [45].

Our device is fabricated using electron-beam lithography (EBL) to pattern components onto a 100 nm aluminum ground plane. The ground plane is connected across coplanar waveguides through crossovers, realized as aluminum air bridges. The Josephson junctions are aluminum based, fabricated using EBL followed by double-angle evaporation. The junction electrodes are shorted to the base layer metallization using Al-based bandages.

TABLE I. Summary of device parameters. Transmon frequencies and coherence times, coupling strengths g_{ij} , and readout resonator (RO) frequencies ω_{RO} and κ, χ values, reported at the qubit idling points. Coherence values are the median of ~ 25 measurements, with the listed uncertainty being the standard deviation of those measurements (* except T_2^* for Q3, which is from a single measurement). All coherence fits are to exponential decays which properly capture the qualitative behavior but neglect some small systematic deviations from exponential behavior such as beating in Ramsey T_2^* fringes. The coupling g_{12} is directly measured, while the couplings g_{13} and g_{23} are calculated based on the measured ancilla dispersive shifts in Appendix G.

Property	Q1	Q2	Q3 (ancilla)
$\omega_{\text{min}}/2\pi$ (GHz)	3.1	3.3	2.5
$\omega_{\text{max}}/2\pi$ (GHz)	5.1	6.1	3.95
$\omega_{\text{idle}}/2\pi$ (GHz)	5.1	5.1	3.74
$\eta/2\pi$ (MHz)	193	204	196
T_1 (μs)	36(12)	14(4)	38(5)
T_2^* (μs)	31(14)	1.29(6)	4.4(2)*
$g_{12}/2\pi$ (MHz)	90.1	90.1	
$g_{13}/2\pi$ (MHz)	8.4		8.4
$g_{23}/2\pi$ (MHz)		81.7	81.7
$\omega_{\text{RO}}/2\pi$ (GHz)	7.749	7.511	7.341
$\chi_{\text{RO}}/2\pi$ (MHz)	3.73(7)	0.32(1)	2.53(3)
$\kappa_{\text{RO}}/2\pi$ (MHz)	9.3(4)	0.87(2)	6.7(2)

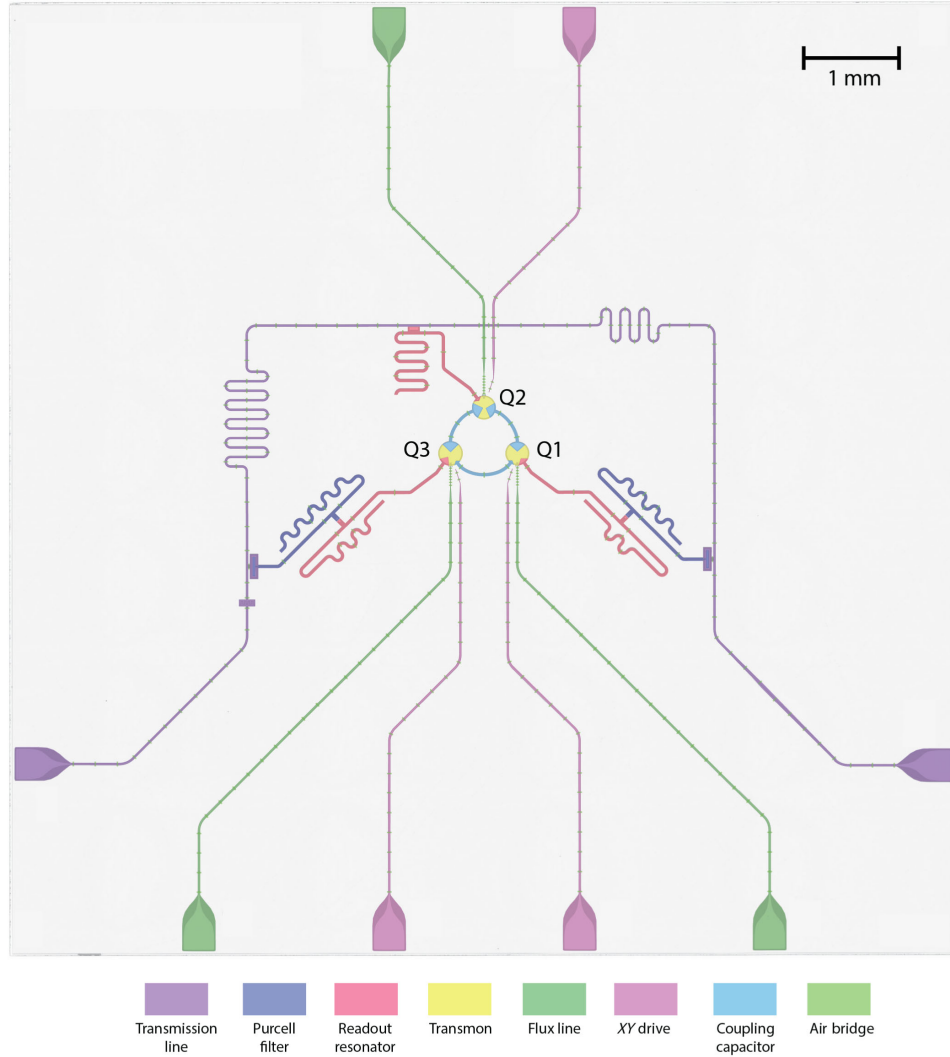


FIG. 6. Optical image of superconducting device. False-color overlays indicate the functional purpose of each component needed to define the transmon qubits, control them with microwave and flux lines, and perform readout.

The device is wire bonded to a printed circuit board and cooled to ~ 10 mK at the base of a Bluefors dilution refrigerator. Microwave and baseband signals are delivered from custom control hardware based on a Xilinx RFSoc, with the signal chains shown in Fig. 7.

APPENDIX B: TRANSMON READOUT

Fast transmon readout, particularly for the ancilla qubit, is critical for erasure qubit experiments. We use a traveling wave parametric amplifier (TWPA) [46] from MIT Lincoln Laboratory to amplify readout signals and Purcell filters to facilitate fast readout without degrading qubit lifetimes. The ancilla readout consists of a 200 ns microwave tone on the transmission line, while we integrate the readout signal against a linear matched filter which extends for an additional 140 ns after the drive tone ends [47]. We assign thresholds in the IQ plane to preferentially minimize

false-positive rates (typically reaching $\lesssim 0.1\%$) at the cost of slightly increasing false-negative rates, which are mainly limited by the few-percent probability of T_1 decay during readout. This choice of threshold is helpful to minimize false-positive rates during erasure detection.

The dual-rail transmons Q1 and Q2 are read out at the end of the circuit over $1\ \mu\text{s}$. During this period, Q2 is flux pulsed up to its maximum frequency at 6.1 GHz, while Q1 remains at 5.1 GHz. For the data in Fig. 5 in which we tune the operating point of the dual-rail qubit below 5.1 GHz, the final readout always occurs while Q2 is flux pulsed back to its maximum frequency.

We observe that for some operating points, reading out the ancilla can stimulate transitions on the dual-rail qubit. We attribute these to measurement-induced state transitions [48,49]. In particular, with the ancilla idling at 3.74 GHz, we find that if the ancilla is excited, there is a small, 0.0293(3)%, chance that its readout will excite the dual-rail

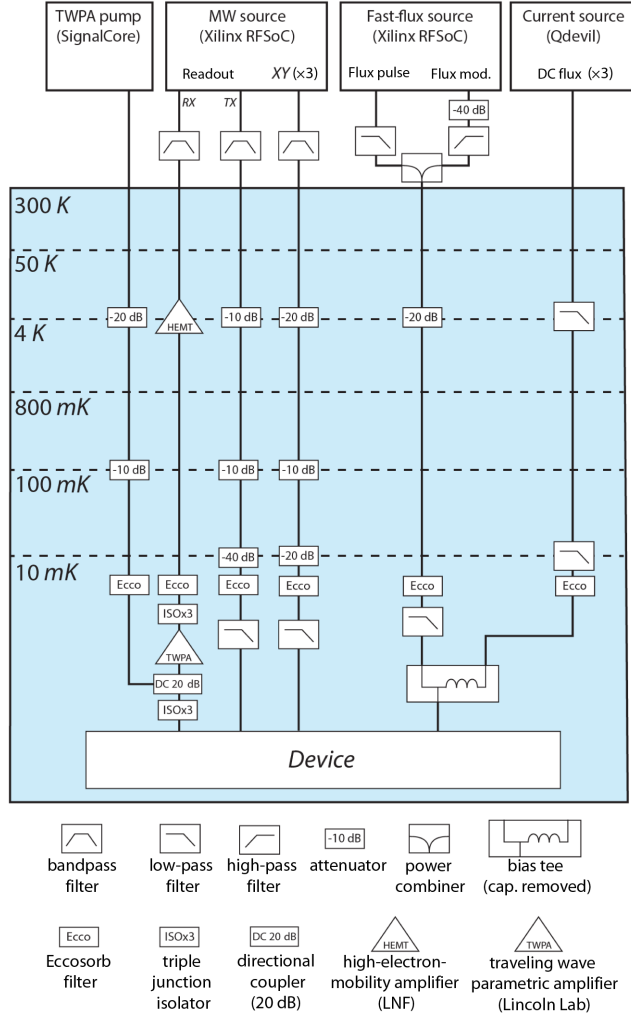


FIG. 7. Controls and cryogenic wiring. We show the signal chain including components at various temperature stages within the dilution refrigerator. A single readout chain is used for the processor. Each qubit has an XY control line, as well as a slow-flux line which is combined with a fast-flux line on a bias tee with the capacitor removed from the fast-flux side to enable low-frequency flux pulses. The fast-flux line for Q2 is generated by combining a signal which is low-pass filtered (cutoff < 22 MHz), used for flux-assisted readout, along with another source which is attenuated and high-pass filtered (cutoff > 41 MHz), used for flux modulation at $2g = 2\pi \times 180$ MHz to drive single-qubit gates; this setup mitigates high-frequency control noise which could limit the dual-rail T_1 . The fast-flux lines for the other qubits (Q1 and Q3) are not needed for these experiments and are terminated with 50Ω at room temperature (not shown).

pair from $|00\rangle$. This only occurs if the ancilla was actually excited, which already flags that an erasure error occurred on the dual-rail pair, and thus postselection against erasure handles these errors. We further note that other types of measurement-induced transitions, including fluctuating behavior and transitions when the ancilla is in $|0\rangle$, were observed at other operating points for the ancilla, including at its maximum frequency of 3.95 GHz, motivating the

operation of the ancilla below its maximum in this work. In the future, such excitation events may be avoided by alternate allocation of frequencies, or can be addressed by resetting both the ancilla and data qubits after detection of an erasure event.

APPENDIX C: DUAL-RAIL HAMILTONIAN AND CONTROLS

The Hamiltonian for the three transmons comprising the dual-rail and ancilla system is given as follows, treating the transmons as nonlinear oscillators:

$$H = \sum_{i=1}^3 \left(\omega_i a_i^\dagger a_i - \frac{\eta_i}{2} a_i^\dagger a_i^\dagger a_i a_i \right) + \sum_{i<j} g_{ij} (a_i^\dagger a_j + a_j^\dagger a_i), \quad (\text{C1})$$

with the values of these parameters provided in Table I. We work in a regime where the detuning between the ancilla and the transmons, $\Delta = \omega_3 - \omega_2 \simeq \omega_3 - \omega_1$, is large compared with the anharmonicity $\eta = \eta_1 \approx \eta_2 \approx \eta_3$ and the intertransmon couplings g_{ij} . As such, the system is well described by an effective dispersive Hamiltonian coupling the ancilla and the dual-rail qubit.

This can be derived using the standard blackbox quantization approach: One first diagonalizes the quadratic parts of the above Hamiltonian, and writes the nonlinearity in terms of the resulting eigenmodes (see, for example, Ref. [14]). Keeping only resonant terms, and projecting to the relevant states of the dual-rail qubit ($|00\rangle$, $|0_L\rangle$, $|1_L\rangle$), we have the following effective (rotating-frame) Hamiltonian:

$$H = \frac{1}{2} (\chi_0 |0_L\rangle\langle 0_L| + \chi_1 |1_L\rangle\langle 1_L|) \sigma_z^{\text{anc}}, \quad (\text{C2})$$

where $\chi_0 \approx \chi_1 \approx 2\eta(g_{23}^2/\Delta^2) \approx 2\pi \times 1.5$ MHz (see Appendix G). This amounts to half the usual dispersive coupling between a pair of transmons (Q2 and Q3), which can be understood from the fact that each dual-rail mode has half its weight in Q2.

The difference $\delta\chi = \chi_1 - \chi_0$ arises due to subleading terms in g_{12}/Δ (which lead to a small difference in detuning between the ancilla and the two dual-rail modes), as well as the presence of the small additional coupling g_{13} and potentially a detuning $\delta = \omega_1 - \omega_2$. One obtains

$$\frac{\delta\chi}{\bar{\chi}} = 2 \frac{g_{12}}{\Delta} + 2 \frac{g_{13}}{g_{23}} - \frac{\delta}{g_{12}}, \quad (\text{C3})$$

to first order in η/Δ , g_{ij}/Δ , g_{13}/g_{23} , and δ/g_{12} , where $\bar{\chi}$ is the average of χ_0 and χ_1 . The nonresonant terms in the expansion of the nonlinear Hamiltonian add an additional perturbative contribution to $\delta\chi/\bar{\chi}$ which is of second order in the above small parameters. In our parameter regime,

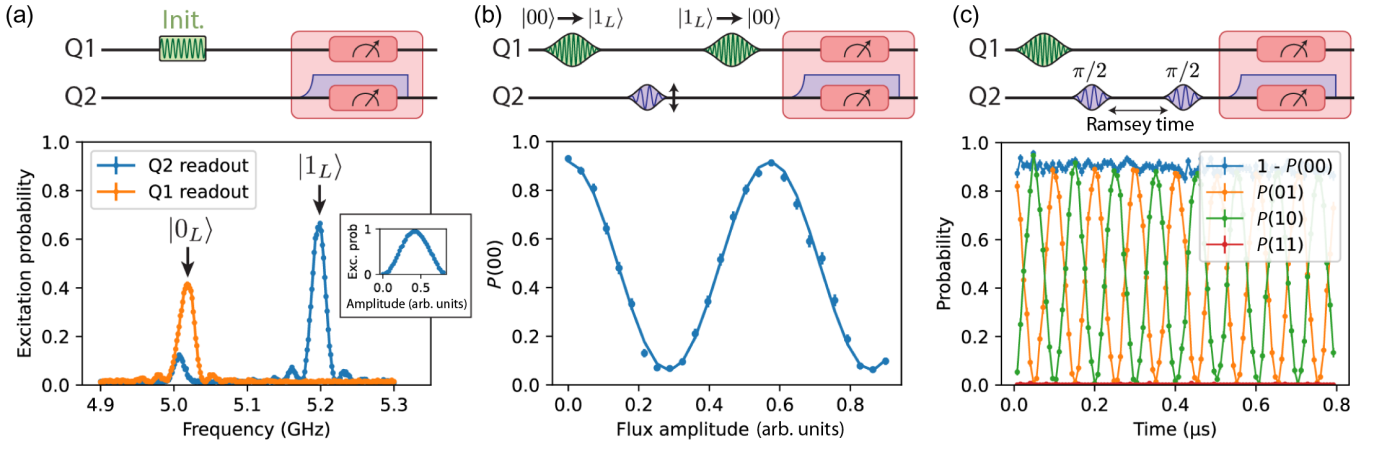


FIG. 8. Dual-rail calibration procedure. (a) After parking the transmons on resonance, we perform microwave spectroscopy to identify the two hybrid modes $|0_L\rangle$ and $|1_L\rangle$ at $\omega_0 \pm g$. A small additional feature at $\omega_0 - \eta/2$ is observed, which is a two-photon resonance going to higher excited states; its near overlap with the transition to $|0_L\rangle$ is particular to these parameters where $g \sim \eta/2$, so for this device we choose to initialize $|1_L\rangle$. Inset: excitation probability as we scan the pulse amplitude on the $|1_L\rangle$ resonance. (b) We calibrate flux-modulation gates by applying a flux-modulation pulse of variable amplitude in between two microwave pulses; this measurement detects population transfer from $|1_L\rangle$ to $|0_L\rangle$ induced by the flux pulse, and is used to approximately calibrate a $\pi/2$ or π pulse within the dual-rail subspace. (c) A Ramsey experiment used to measure the dual-rail frequency relative to a reference frequency. We show here the full bit-string probabilities of the Q1 and Q2 readout; the 01 and 10 results are interpreted as representing population in $|0_L\rangle$ and $|1_L\rangle$, respectively, prior to the adiabatic separation of qubits during readout.

$|\delta\chi/\bar{\chi}| \ll 1$, enabling the erasure-check approach illustrated in Fig. 1(e).

Control fields. The two main control fields for the dual-rail qubit are (1) microwave initialization ($|00\rangle \rightarrow |1_L\rangle$) and (2) flux modulation for driving single-qubit gates. Since the dual-rail modes are hybridized across Q1 and Q2, these fields can in principle be applied to either transmon. In particular, a charge drive applied to either transmon can couple from $|00\rangle$ to both logical states [as is visible in the spectroscopy signal in Fig. 8(a)].

For single-qubit gates, we consider a simplified Hamiltonian for Q1 and Q2, given by $H = g(\sigma_1^+ \sigma_2^- + \text{H.c.}) + \delta(t)\sigma_2^z$ describing a pair of coupled qubits with a time-dependent differential frequency [28]. Modulating $\delta(t)$ at frequency $2g$ drives resonant rotations within the logical subspace. This frequency modulation can be realized by flux modulation on either transmon, but it is preferable to modulate the transmon whose frequency versus flux sensitivity is most linear to avoid shifts in the average frequency during modulation.

APPENDIX D: DUAL-RAIL QUBIT CALIBRATION

The dual-rail qubit calibration procedure is as follows.

- (1) Identify approximate operating points. Spectroscopically measure the avoided crossing between the two transmons to identify which flux offsets bring the transmons on resonance [as shown in Fig. 1(b)].

- (2) Initialize the dual-rail qubit. With both transmons parked at the operating point, apply a microwave drive at one of the two hybrid mode frequencies and scan its amplitude to calibrate a π pulse from $|00\rangle$ to $|0_L\rangle$ or $|1_L\rangle$ [Fig. 8(a)].
- (3) Calibrate gates within the subspace. To calibrate an approximate $\pi/2$ and π pulse induced by flux modulation of one qubit, we perform two microwave π pulses on the $|00\rangle$ to $|1_L\rangle$ transition with the flux-modulation pulse applied in between to detect population transfer from $|1_L\rangle$ to $|0_L\rangle$. We scan the amplitude of the flux-modulation pulse to obtain an approximate $\pi/2$ and π pulse [Fig. 8(b)].
- (4) Refine flux bias on transmons. We subsequently refine the calibration by performing a Ramsey sequence on the dual-rail qubit to more precisely measure the dual-rail frequency [Fig. 8(c)]. We repeat this as a function of flux offset on one qubit to find the operating point which minimizes the dual-rail frequency, which optimizes the robustness of the dual-rail qubit to flux noise (example shown in Fig. 15).

APPENDIX E: REFINING SINGLE-QUBIT GATE CALIBRATION

The single-qubit gate is a flux-modulation pulse on Q2 at the dual-rail frequency $2g = 2\pi \times 180$ MHz. The pulse is 48 ns with a Gaussian envelope ($\sigma = 12$ ns). We calibrate the amplitude of the pulse by repeated application to optimize the rotation angle [50]. We subsequently calibrate

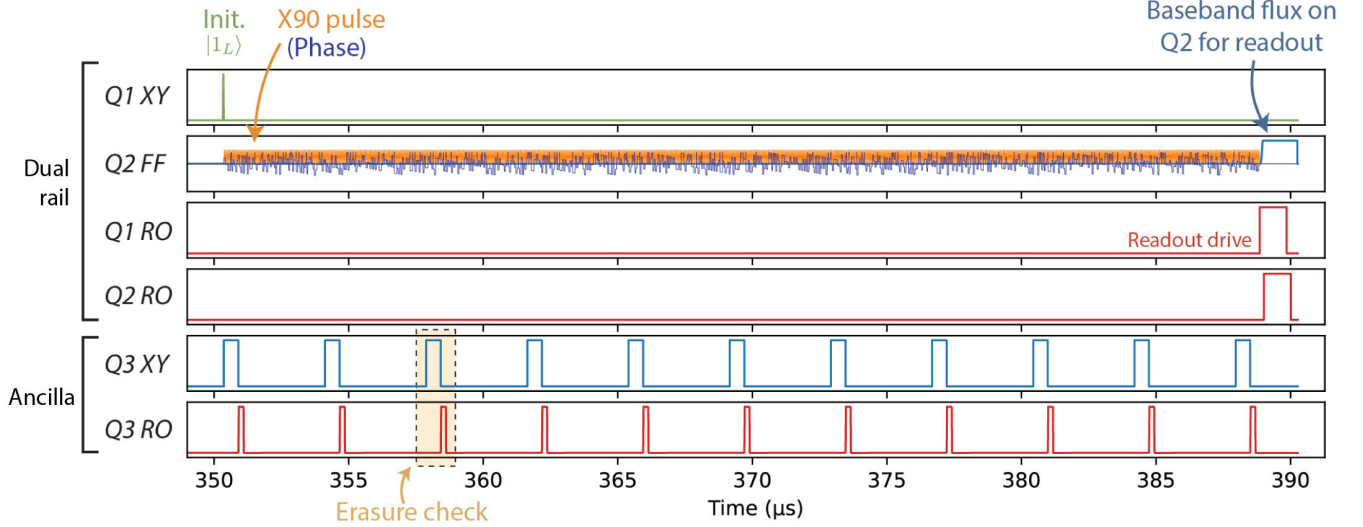


FIG. 9. Experimental sequence for randomized benchmarking. The sequence begins with $\sim 350 \mu\text{s}$ idle time for the transmons to relax into $|0\rangle$. A microwave pulse on Q1 initializes the dual-rail qubit in $|1_L\rangle$. A sequence of Cliffords is then applied, each composed of two X90 pulses with a particular phase shift of the drive field. Shown here is a sequence with 400 Cliffords (800 X90s). In parallel with this sequence, 11 erasure checks are performed, evenly spaced across the full circuit. Finally, a baseband flux pulse is applied to Q2 to adiabatically separate Q1 and Q2, during which these qubits are read out. FF stands for fast flux.

a $Z(\phi)$ phase correction with $\phi = 0.067$ rad to be applied after every X90 due to small shifts in the dual-rail frequency during the gate [33,51]. Z rotations are applied virtually by shifting the phase of the flux-modulation rf drive [33,52]. The speed of the gate can be optimized to approach the Larmor period of ~ 6 ns set by the inverse of the qubit frequency $2g = 2\pi \times 180$ MHz, enabling similar gate speeds to that of transmons with typical $2\pi \times 200$ MHz anharmonicity.

Other single-qubit gate schemes are also possible with rf qubits such as the dual-rail qubit, closely analogous to single-qubit gate schemes with heavy fluxonium [53]. In particular, baseband flux pulses have been demonstrated to realize exact gates at the maximum speed given by the Larmor period of the qubit [28]. Such fast gates may be optimal for improving single-qubit gate fidelity, but require more complex calibration and preclude the use of phase shifts on the rf field to realize virtual Z gates.

APPENDIX F: RANDOMIZED BENCHMARKING DETAILS AND SUPPLEMENTARY DATA

We show in Fig. 9 an example timing diagram for a single-qubit randomized benchmarking sequence with 400 Cliffords. This sequence is representative of nearly all experiments presented in this work, with the main variability being the sequence of flux-modulation pulses being applied on the fast-flux line of Q2 to induce a particular choice and timing of single-qubit gates.

In the main text, the erasure error and residual error per gate are computed using postselection based on both the midcircuit erasure checks and the final dual-rail readout.

This evaluates most directly the gate fidelity within the dual-rail subspace, separating out the performance of midcircuit erasure detection and the final readout. If we instead postselect only on the midcircuit erasure checks, we find nearly the same gate error rates as shown in Fig. 10.

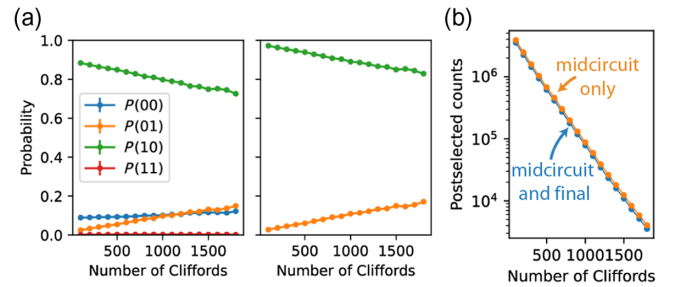


FIG. 10. Randomized benchmarking postselection protocols. (a) The final dual-rail readout, postselected only on midcircuit erasure checks (left-hand panel), shows a gradual decline in the ideal outcome $P(10)$, as well as a background of $P(00) \approx 10\%$ due to readout errors. The $P(10)$ trace does not decay to a known offset, but a fit to Ae^{-pN} is valid for short times and gives an error of $p = 1.10(2) \times 10^{-4}$ per Clifford. This result is close to the error rate evaluated with additional postselection against $|00\rangle$ and $|11\rangle$ in the final readout (right-hand panel), which as reported in the main text has error $1.01(1) \times 10^{-4}$. (b) The erasure error per Clifford is computed based on the decay of postselection probability with increasing circuit depth, and gives nearly the same fitted erasure probability per Clifford when using just midcircuit erasure checks [0.437(3)%] as when using both midcircuit checks and final readout as reported in the main text [0.438(3)%].

APPENDIX G: DISPERSIVE SHIFTS OF THE DUAL-RAIL QUBIT ON THE ANCILLA

The ancilla qubit's frequency is shifted depending on whether the dual-rail pair is in $|00\rangle$, $|0_L\rangle$, or $|1_L\rangle$. While we illustrate this shift with microwave spectroscopy in Fig. 1(e), we also use a Ramsey sequence to precisely characterize these shifts. Specifically, we perform a Ramsey experiment on the ancilla to measure its frequency after initializing the dual-rail qubit in its three possible states. We find that the dispersive shift on the ancilla when the dual-rail goes from $|00\rangle$ to $|0_L\rangle$ and $|1_L\rangle$ is $2\pi \times [1.514(4), 1.610(4)]$ MHz, respectively. The slight differential shift between the two logical states is $2\pi \times 96(4)$ kHz.

APPENDIX H: DETERMINISTIC PHASE SHIFT INDUCED BY ERASURE CHECK

The ancilla-based erasure check induces a small deterministic phase shift on the dual-rail qubit, in addition to a potential dephasing error discussed in Appendix I. To characterize the deterministic phase shift, we repeat the experiment performed in Fig. 4(c) in which we insert a variable number of erasure checks into a fixed-time spin echo experiment; rather than inserting them in a balanced way (equal number of checks in both halves of the spin-echo sequence), we insert all checks in the first half or all checks in the second half (Fig. 11). By comparing the final dual-rail phase (measured by scanning the phase of the final $\pi/2$ pulse), we assess that the induced phase shift is

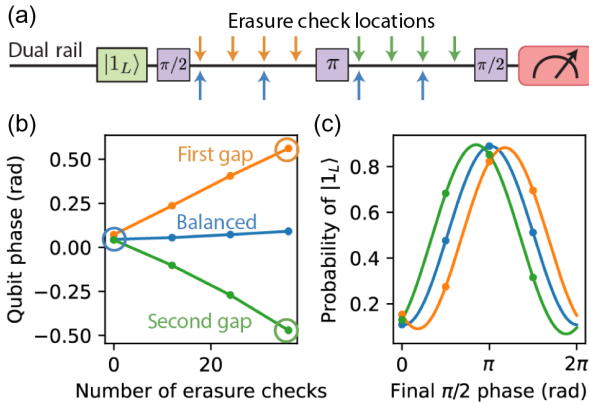


FIG. 11. Deterministic phase shift induced by erasure check. (a) In a similar experiment to that presented in Fig. 4, we perform a spin echo and insert a variable number of erasure checks into either the first half, the second half, or both halves of the sequence. (b) We see that the phase of the dual-rail qubit is essentially independent of the number of erasure checks when balanced across both arms, but picks up a phase shift proportional to the number of checks when all checks are applied within one arm. The phase accumulation per check fits to $0.0142(3)$ rad. (c) We illustrate example phase measurements of the points circled in (b), in which we scan the phase of the final $\pi/2$ pulse in the spin-echo sequence and fit to a sinusoidal model.

$0.0142(3)$ rad per check. This deterministic shift is correctable in circuits with virtual Z gates.

APPENDIX I: BOUNDING DEPHASING ERROR INDUCED BY ERASURE CHECKS

In Fig. 4(c) of the main text, we present a spin-echo experiment in which we measure the remaining phase coherence of the dual-rail qubit after $114 \mu\text{s}$ as a function of how many erasure checks are inserted in the gap. As we add more erasure checks, we eliminate shots for which the system decayed to $|00\rangle$ and heated back into the subspace. This effect should plateau once there is a certain density of erasure checks, which we expect to be the plateau observed starting at $N \geq 16$ checks. Assuming this model, the residual decay from 16 to 128 checks gives a dephasing error per check of $p_{\text{err}} = 0.035\%$.

This analysis, however, is not rigorous as we cannot rule out that as we apply more postselection, the coherence should appear to continue to grow, and that this effect may be balanced by a larger amount of dephasing from more checks. We can bound such a model in the following way. To start, we evaluate the worst-case scenario, in which the *true* coherence is perfect after the $114 \mu\text{s}$ evolution time, and that the measured phase coherence of $77.4(1.1)\%$ after 128 checks is purely due to dephasing from the checks: $0.774 \geq e^{-Np_{\text{err}}}$, from which we conclude a rigorous upper bound on the error $p_{\text{err}} \leq 0.2\%$.

This worst-case scenario is unrealistic, as it requires (1) perfect idling coherence of the dual-rail qubit in the absence of erasures and (2) the low coherence after $N = 16$ erasure checks is fully explained by the imperfect postselection. We improve these assumptions as follows: To begin, we bound how much imperfect postselection can contribute at $N \geq 16$ based on independent T_2 coherence measurements where we use 11 checks and more dynamical decoupling. In such experiments, we have $\geq 85\%$ coherence remaining after $114 \mu\text{s}$, including readout and pulse errors; this means that the imperfect postselection can only contribute $\leq 15\%$ of the loss of coherence for $N = 16$; such a bound implies a slightly stronger upper bound for the erasure-check dephasing error of $p_{\text{err}} \leq 0.16\%$. Even so, this bound assumes that in the absence of erasures, the dual rail retains 95% coherence after $114 \mu\text{s}$, consistent with an exponential decay time of $T_2^{\text{echo}} = 2.2$ ms. For a general T_2^{echo} , we would have that $0.774 = e^{-\tau/T_2^{\text{echo}}} e^{-Np_{\text{err}}}$, or equivalently, $p_{\text{err}} = [-\ln(0.774) - \tau/T_2^{\text{echo}}]/128$. Plugging in our independently measured $T_2^{\text{echo}} = 540 \mu\text{s}$ presented in Fig. 2 which predicts that the true coherence remaining would only be $\sim 80\%$, consistent with the observed plateau, this would give us a smaller error estimate of $p_{\text{err}} \approx 0.035\%$. We hypothesize that this error may be related to measurement-induced state transitions while reading out the ancilla. In this work, we report a more conservative upper bound of $p_{\text{err}} < 0.1\%$, which could be violated only if

$T_2^{\text{echo}} \gtrsim 900 \mu\text{s}$, and leave a more focused investigation for future work.

APPENDIX J: LIMITS IMPOSED BY TRANSMON HEATING

The primary leakage mechanism of the dual-rail qubit is T_1 decay of underlying transmons that leaves the system in $|00\rangle$. However, the transmons individually have a slow heating timescale which is given by $T_{\text{heat}} \approx T_1/p_{\text{equil}}$, where p_{equil} is the thermal excitation probability of the individual transmons. We measure $p_{\text{equil}} = 0.2(1)\%$ (using single-shot readout methods), so while the T_1 decay which leads to the erasure lifetime of the dual rail is $\sim 30 \mu\text{s}$, the heating timescale is $\sim 15 \text{ ms}$.

This heating mechanism gives rise to two important effects. Firstly, population that has accumulated in $|00\rangle$ can return into the dual-rail subspace through heating. If we only postselect against leakage based on the final readout of the dual-rail pair, then such decay and reheating events, in which all phase coherence was lost, are not detectable and appear to limit the dual-rail coherence. Surprisingly, despite the long $> 10 \text{ ms}$ timescale for heating, the timescale for dephasing due to decay and reheating is actually the much shorter time T_1 .

To see why, we compare the coherent population which remains in the dual-rail subspace (without decaying) to the total population ending in the dual-rail subspace including decay and reheating. The probability of never decaying is e^{-t/T_1} , while the probability of ending in the dual-rail subspace is given by thermal equilibration dynamics: $(1 - 2p_{\text{equil}})e^{-t/T_1} + 2p_{\text{equil}}$, where $2p_{\text{equil}}$ is double the single-transmon thermal population as there are two transmons which may be excited. The ratio between these defines the coherence function:

$$C(t) = \frac{e^{-t/T_1}}{(1 - 2p_{\text{equil}})e^{-t/T_1} + 2p_{\text{equil}}}. \quad (\text{J1})$$

The dephasing time where $C(t) = 1/e$ is

$$T_{\text{deph}} \approx T_1 \ln \frac{e - 1}{2p_{\text{equil}}}, \quad (\text{J2})$$

which is several times T_1 for typical values of p_{equil} (i.e., $T_{\text{deph}} \approx 6T_1$ for $p_{\text{equil}} = 0.2\%$), with only a logarithmic dependence on p_{equil} or equivalently the heating timescale T_{heat} .

Midcircuit erasure detection, however, effectively mitigates the contribution of this effect as long as the erasure checks are performed sufficiently frequently. Analytically, for N perfect erasure checks evenly spaced in a total time t , the remaining coherence would be $C(t) \rightarrow C(t/N)^N$. Experimentally, we compare dual-rail coherence measurements with varying numbers of erasure checks in Fig. 12 to

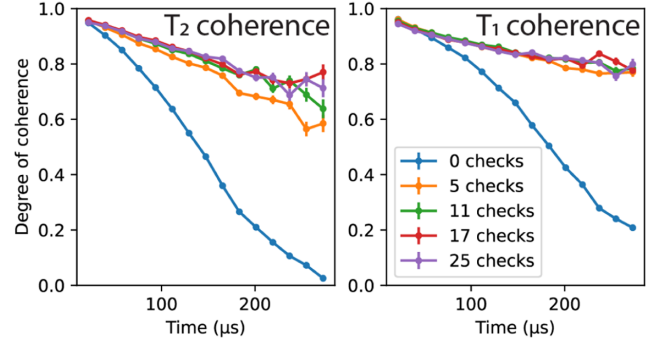


FIG. 12. Varying the number of erasure checks during coherence measurements. When measuring dual-rail T_2 or T_1 with no erasure checks, the coherence appears to rapidly decay over 200–300 μs due to leakage to $|00\rangle$ which then heats back into the dual-rail subspace. As we add more erasure checks (evenly spaced throughout the sequence for each variable length), we find that 11 erasure checks is enough to eliminate the contribution of such instances so that we can more accurately probe coherence within the subspace.

ensure that we mitigate this contribution and are not limited by this mechanism after $\sim 300 \mu\text{s}$ of evolution. In the main text, we use 11 erasure checks, evenly spaced during the free evolution time, for the data in all figures.

The second leakage mechanism is direct heating from the dual-rail subspace into the two-photon subspace spanned by $|11\rangle$, $|02\rangle$, and $|20\rangle$, which, due to the additional channels for heating, has a $3\times$ combined rate compared to the heating rate $\Gamma_{0 \rightarrow 1}$ of a single transmon (i.e., from $|01\rangle$ the heating into $|11\rangle$ occurs at rate $\Gamma_{0 \rightarrow 1}$ while the heating into $|02\rangle$ occurs at $2\Gamma_{0 \rightarrow 1}$ due to the larger matrix element). Such heating events will be followed by decay back into the dual-rail subspace and will similarly appear as decoherence with timescale $T_2 \sim T_{\text{heat}}/3 \sim 5 \text{ ms}$. This will typically set an upper limit on how much coherence can be observed within the dual-rail subspace, as such events are not caught by midcircuit erasure detection as currently implemented. To circumvent this limit, one could expand the erasure detection protocol to detect not only leakage to $|00\rangle$ but also leakage to the two-photon subspace. Such an approach is a feasible extension of the ancilla qubit method described in this work, and may reduce the requirements on transmon temperature to enable working at higher cryostat temperatures.

APPENDIX K: TELEGRAPH NOISE IN THE DUAL-RAIL QUBIT

The main text presents dual-rail coherence with dynamical decoupling. Without dynamical decoupling, the dual-rail Ramsey coherence exhibits fluctuating behavior, sometimes showing “beat-note-like” decay and revival [see Fig. 13(a) for two representative Ramsey measurements averaged over $\sim 1 \text{ min}$].

By instead using short Ramsey experiments to measure the dual-rail frequency in second-scale intervals, we

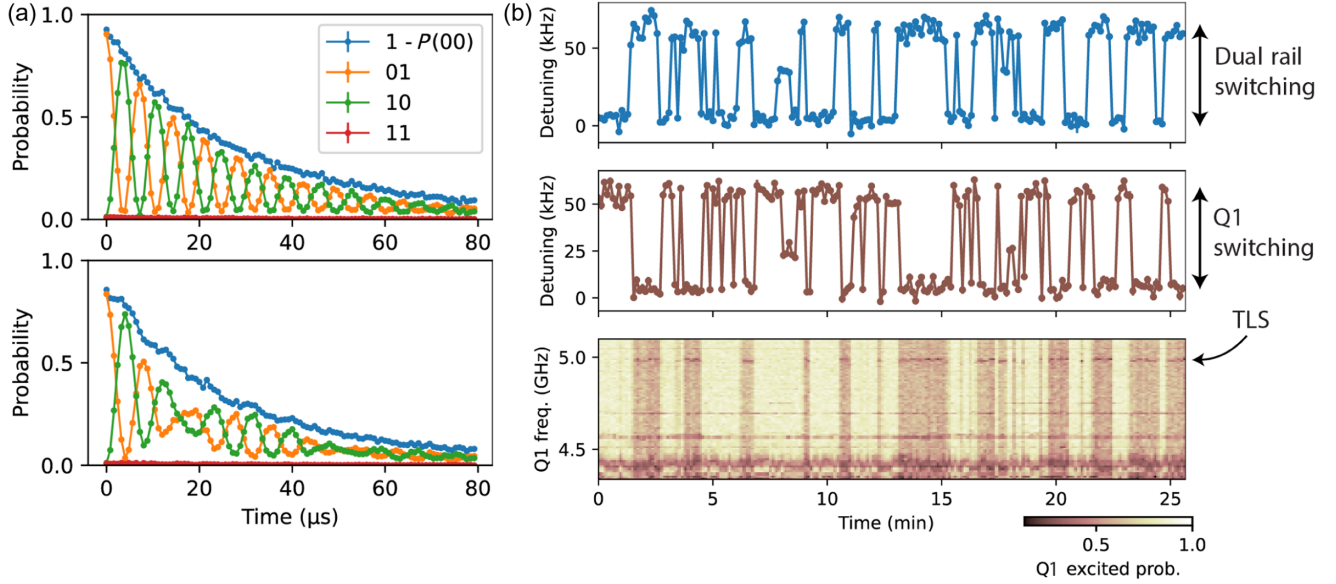


FIG. 13. Telegraph noise on the dual-rail qubit. (a) We show two typical Ramsey T_2^* measurements, integrated for ~ 1 min. Within the envelope of not decaying to $|00\rangle$, we see fringes which track the phase coherence of the dual-rail qubit. We find in some cases stable fringes (upper plot) and in other cases decay-revival behavior (lower plot). (b) By repeatedly measuring the dual-rail frequency over shorter time intervals, we identify telegraphlike switching in the frequency. Integrating a Ramsey experiment for longer than this switching time gives rise to the beating effect seen in (a). By interleaving this time-resolved measurement with similar time-resolved measurements of Q1's frequency and flux-pulse spectroscopy measurements on Q1, we note that this switching is present in similar scales on both the dual-rail and Q1 and is tied to the toggling of a TLS-like mode at 4.98 GHz. When the TLS is present, it reduces the readout fidelity of Q1 leading to a reduced contrast in the spectroscopy slice.

directly observe that this behavior is explained by telegraph noise, with a switching amplitude of ~ 60 kHz and switching timescale of tens of seconds [Fig. 13(b), upper panel]. Interleaving such measurements with Ramsey measurements on Q1 when parked alone at the same operating point, we find that both exhibit similar scale and time-correlated switching [Fig. 13(b), middle panel].

We attribute this to a flickering TLS mode at 4.98 GHz by interleaving the above frequency measurements with a spectroscopy measurement on Q1 in which we initialize Q1 in $|1\rangle$ and flux pulse it to a variable frequency for $1 \mu\text{s}$ and observe if the excitation is lost [45]. We find a clear correlation between the presence of a dark mode at 4.98 GHz with the frequency shift observed both on Q1 alone and on the dual-rail qubit.

While this telegraph noise is sufficiently slow that dynamical decoupling easily addresses it, interactions between the dual-rail qubit and proximal TLS modes warrant further study to understand what types of impacts TLSs may have on dual-rail qubits beyond just dispersive shifts and also potentially to probe TLS physics with a complementary toolbox to transmon probes (see Appendix N 3 for further discussion).

APPENDIX L: DUAL-RAIL QUBIT STABILITY OVER TIME

While the dual-rail frequency flickers on slow timescales due to telegraph noise, the behavior with even minimal

dynamical decoupling (i.e., a single echo pulse) shows much more stable performance over long time intervals. The coherence data presented in the main text Fig. 2 was averaged over ~ 24 hours, and in Fig. 14 we show the data in discrete chunks during this window.

From the same datasets, we similarly extract the erasure lifetime over the measurement period, and find moderate fluctuations between 20 and $40 \mu\text{s}$, consistent with typical T_1 fluctuations for individual transmons. Somewhat surprisingly, we note that the two dual-rail logical states exhibit different erasure lifetimes, each of which fluctuate seemingly independently. While naively we would assume that each mode would have the same erasure lifetime given by the average T_1 decay rate of the two constituent transmons, this differential may be due to frequency-dependent effects which are differently resonant with the two hybrid modes.

APPENDIX M: DUAL-RAIL COHERENCE AND NOISE SUPPRESSION

The energy gap of the dual-rail qubit is determined primarily by the capacitive coupling strength g between transmons, and is only weakly sensitive to the detuning between the transmons. For transmon frequencies ω_1, ω_2 , the dual-rail energy is given by

$$E_{\text{DR}} = \sqrt{(2g)^2 + (\omega_1 - \omega_2)^2} \approx 2g + \frac{(\omega_1 - \omega_2)^2}{4g}, \quad (\text{M1})$$

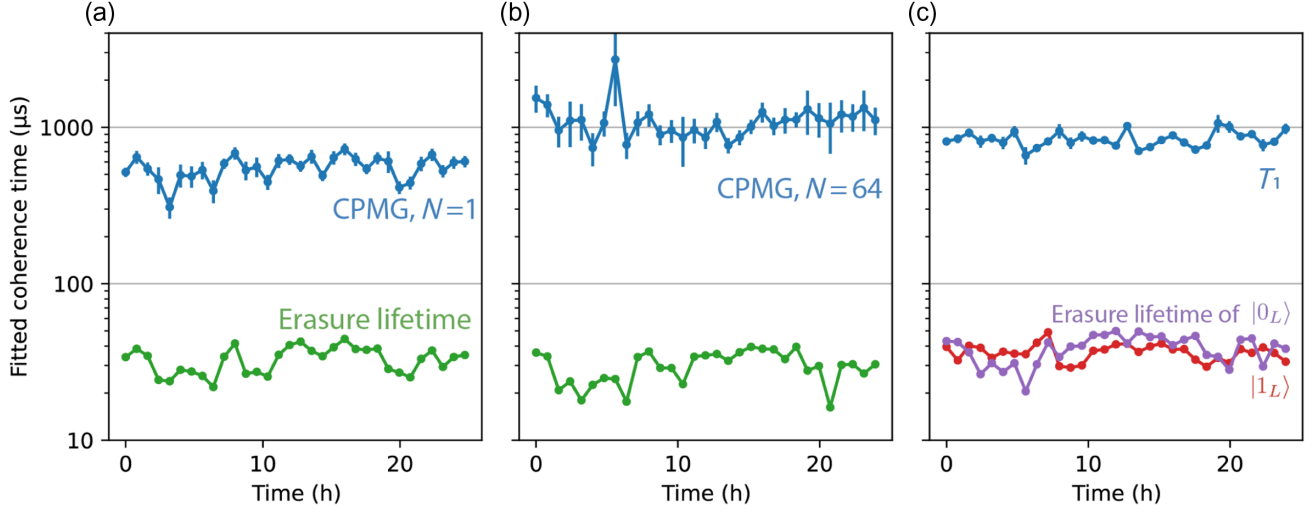


FIG. 14. Stability of dual-rail coherence metrics. The coherence measurements presented in Fig. 2 were averaged over data acquired over ~ 24 h. Here, we analyze subsets of the same dataset in a time-resolved way to track the fitted coherence metrics during that period. We show three representative metrics: (a) CPMG with one π pulse, (b) CPMG with 64 π pulses, and (c) T_1 within the dual-rail subspace. In each plot, we additionally show the erasure lifetime extracted during each measurement. For the T_1 measurement, this erasure lifetime is extracted separately for the two logical states and found to be fluctuating differently for the two states, despite the measurements being taken in an interleaved fashion. Large discrepancies between the erasure lifetimes of the two logical states can lead to dephasing of the dual-rail qubit, but this effect is mitigated by dynamical decoupling.

thus offering quadratic suppression of the detuning between transmons. However, the Hamiltonian coupling strength g , produced by a fixed capacitance between the qubits, itself scales with qubit frequency as $g \propto \sqrt{\omega_1 \omega_2}$. Defining a reference operating frequency for the qubits ω_0 with corresponding coupling strength g_0 , and with relative detunings $\delta_{1,2} = \omega_{1,2} - \omega_0$, the energy gap takes the following form assuming $\delta_{1,2} \ll g_0 \ll \omega_0$:

$$E_{\text{DR}} = 2g_0 + \frac{g_0}{\omega_0}(\delta_1 + \delta_2) + \frac{(\delta_1 - \delta_2)^2}{4g_0}. \quad (\text{M2})$$

When operating at $\delta_1 = \delta_2 = 0$, the final term offers quadratic suppression, but the overall energy gap retains a linear sensitivity to qubit frequency noise scaled down by $\partial E_{\text{DR}} / \partial \delta_{1,2} = g_0 / \omega_0 \approx 0.018 \approx 1/57$. If the qubits are limited by $1/f$ noise, this would predict that the dual-rail coherence is scaled linearly with this factor.

By introducing a small detuning, however, the dual-rail energy gap can become linearly insensitive to frequency noise on one of the two qubits ($\partial E_{\text{DR}} / \partial \delta_1 = 0$), at the cost of a larger sensitivity to the other ($\partial E_{\text{DR}} / \partial \delta_2 = 2g_0 / \omega_0$). This is done by parking at $\delta_1 = \delta_2 - 2g_0^2 / \omega_0$; i.e., one qubit is parked below the other qubit by an amount $2g_0^2 / \omega_0 \approx 2\pi \times 3.2$ MHz for the current device parameters. Letting ϵ_1, ϵ_2 now be deviations from these target operating points, we have

$$E_{\text{DR}} = \left[2g_0 - \frac{g_0^3}{\omega_0^2} \right] + \frac{2g_0\epsilon_2}{\omega_0} + \frac{(\epsilon_1 - \epsilon_2)^2}{4g_0}. \quad (\text{M3})$$

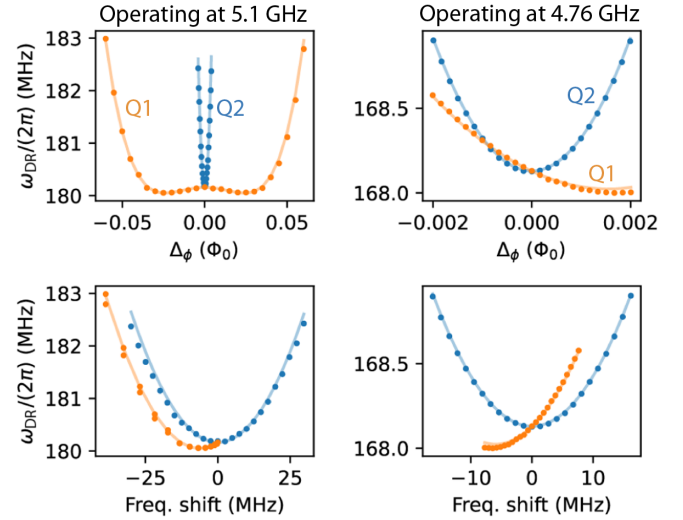


FIG. 15. Measuring dependence of dual-rail frequency on transmon flux. At the two extremal operating points from Fig. 5, we measure how the dual-rail frequency shifts as we shift the flux on one of the two transmons away from its calibrated operating point (upper row). Each point here is measured with a Ramsey experiment on the dual-rail qubit. The functional form is dependent both on how the dual-rail frequency changes with underlying transmon frequency and also how the transmon frequency shifts with flux. By independently calibrating how the transmon frequencies change in this range, we plot the same data with an altered x axis that reflects instead how much the transmon frequency shifted (bottom row). In both rows, the transparent trace is the analytic model from Eq. (M3), with the only adjusted parameters being the calibrated dual-rail energy gap at each operating point.

This is a preferable operating point if one qubit is noisier than the other, which is in practice often the case.

This type of local optimum is found automatically through a calibration procedure in which one transmon is parked at a target operating point (ideally, its flux sweet spot), and then the other transmon is tuned near resonance, where the dual-rail energy gap is measured as a function of flux on the scanning transmon. The operating flux is selected as the one which minimizes the dual-rail energy gap, thus minimizing sensitivity to flux noise. After calibrating in this way, we check the validity of this model by measuring the dual-rail frequency as we scan flux on each qubit and find excellent agreement (Fig. 15).

We use this strategy at each operating point studied in Fig. 5. As a result, we expect the dominant contribution from transmon frequency noise to dual-rail dephasing to come from the linear suppression on Q1. Quantitative estimates of these effects are provided in Appendix N.

APPENDIX N: MECHANISMS FOR DUAL-RAIL DECOHERENCE

Here, we enumerate several mechanisms which can contribute to dual-rail dephasing and where possible estimate the associated timescales.

1. Flux noise

Flux noise induces frequency noise on the individual transmons. The transmon frequency noise can affect the dual-rail qubit in two ways: (1) the dual-rail energy gap is sensitive to frequency noise on the underlying qubits according to Eq. (M3), which can cause dephasing of the dual-rail qubit, and (2) high-frequency noise at the dual-rail energy gap $2g$ can cause T_1 -limiting bit-flip transitions within the dual-rail subspace. The rate of such decoherence mechanisms depends heavily on the noise spectrum; while we do not have a full model for the qubit noise spectrum, we assess the influence of two paradigmatic noise processes below: $1/f$ noise and Johnson-Nyquist noise.

a. $1/f$ noise

To assess the possible influence of $1/f$ noise, we consider a noise spectrum which is scaled in amplitude to give a particular coherence time when applied to a single transmon alone. We focus on the T_ϕ^{echo} dephasing time, that is the pure-dephasing time in a spin-echo experiment on a single transmon, which is preferable to analyze over the Ramsey dephasing time as it is insensitive to the divergence of $1/f$ at low frequencies.

Firstly, we consider bit flips induced by $1/f$ noise extending up to $2g = 2\pi \times 180$ MHz. We consider a two-sided frequency noise spectrum $S(\omega) = 2\pi A^2/|\omega|$ which gives rise to a spin-echo dephasing time for a single transmon of $T_\phi^{\text{echo}} = [A\sqrt{\ln 2}]^{-1}$ [54]. With the same noise

applied to one transmon within a dual-rail pair, the dual-rail T_1 would be limited as $T_1^{\text{DR}} = 2/S(2g)$ [55]; using $A = [T_\phi^{\text{echo}}\sqrt{\ln 2}]^{-1}$, we see that the dual-rail T_1^{DR} is related to the single-transmon coherence as

$$T_1^{\text{DR}} = \frac{2 \ln 2}{\pi} g (T_\phi^{\text{echo}})^2. \quad (\text{N1})$$

For a noise amplitude which gives reasonable qubit coherences of $T_\phi^{\text{echo}} > 3 \mu\text{s}$ (a lower bound for each measured transmon in the dual-rail pair), the limit on the dual-rail coherence would be $T_1^{\text{DR}} > 2.2$ ms. This is likely a conservative upper bound on the true $1/f$ noise contributions, as the noisier qubit Q2's T_2^{echo} is likely limited by extra low-frequency noise from the control hardware which is low-pass filtered to $< 2\pi \times 20$ MHz and would not extend high enough in frequency to reach $2g$; as such, we expect this noise contribution not to limit our measured dual-rail T_1 .

Secondly, we consider the dephasing induced by frequency noise on Q1 which is linearly suppressed by a factor of $2g_0/\omega_0 \approx 1/28$. For $1/f$ noise, the dual-rail coherence would be scaled by this linear factor, with $T_2^{\text{DR}} \approx 28 \times T_\phi^{\text{echo}}$. Q1 exhibits a low pure-dephasing rate at its sweet spot, with $T_\phi^{\text{echo}} > 100 \mu\text{s}$; if this were due to $1/f$ noise, it would correspond to a dual-rail dephasing time of > 2.8 ms, again not limiting for current experiments. In the furthest operating point where Q1 is on its slope, it exhibits a dephasing time of $T_\phi^{\text{echo}} \approx 37(9) \mu\text{s}$, which would correspond to a dual-rail dephasing time of $T_2^{\text{DR}} \approx 1$ ms. While this is not limiting, this additional contribution is consistent with the decrease in T_2^{DR} measured from the best operating point [$T_2^{\text{DR}} \approx 550(30) \mu\text{s}$] to the worst operating point [$T_2^{\text{DR}} \approx 366(20) \mu\text{s}$] as shown in Fig. 16.

Finally, we consider the dephasing induced by frequency noise on Q2 which is quadratically suppressed. We numerically simulate $1/f$ noise, with an amplitude scaled according to the measured $T_2^{\text{Q2}} = 3 \mu\text{s}$, and find that the contributions to dual-rail dephasing are at the level of ~ 35 ms. While this $1/f$ model likely is not an accurate description of Q2's noise spectrum, particularly due the presence of noise from the control hardware, this illustrates the efficacy of noise suppression in the dual-rail qubit, and we leave a further detailed study of noise spectra for future work.

b. Thermal Johnson-Nyquist noise in flux lines

Our flux lines are thermalized at 4 K without significant further attenuation at lower temperature stages. As such, we expect thermal Johnson-Nyquist noise at this characteristic temperature to be carried down the flux lines and affect the qubits. The induced two-sided spectrum describing

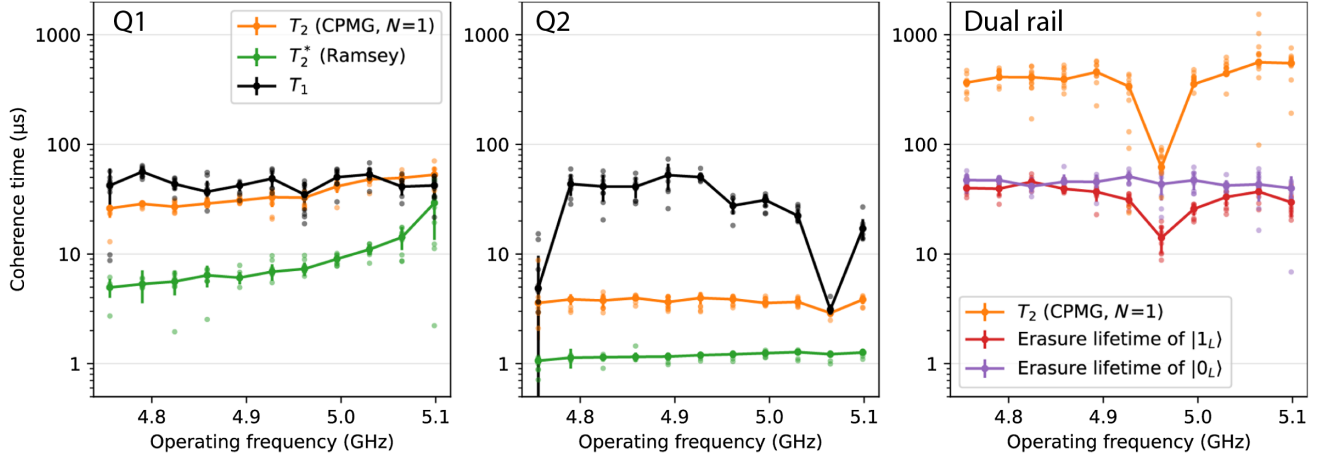


FIG. 16. Supplementary coherence data across tunable operating band. We show T_2^* , T_2^{echo} , and T_1 for both transmons Q1 and Q2 at their range of operating points. For the dual-rail qubit, we show the T_2^{echo} as well as the erasure lifetime for each logical state. We find that the dip in dual-rail coherence at 4.96 GHz corresponds to a drop in the $|1_L\rangle$ erasure lifetime; this upper hybrid mode is at $4.96 \text{ GHz} + 90 \text{ MHz} = 5.05 \text{ GHz}$. We attribute this feature to a TLS coupled to Q2, which is independently seen as the sharp drop in Q2's T_1 around 5.05 GHz. All individual points are fit results for single measurements; for the transmon data for Q1 and Q2, the average result (error bar) is the median of the individual datasets (standard deviation). For the dual-rail T_2 data, the average (error bar) is the fit result of all datasets averaged together (fit uncertainty).

transmon frequency noise due to equilibrium current fluctuations is, for $\omega \ll k_B T$,

$$S(\omega) \approx 2 \left(\frac{\partial \omega_T}{\partial \Phi_e} \right)^2 M^2 \frac{k_B T}{Z}, \quad (\text{N2})$$

where $\partial \omega_T / \partial \Phi_e$ is the sensitivity of the transmon frequency with respect to external flux, M is the mutual inductance to the flux line, k_B is the Boltzmann constant, T is the effective noise temperature, and Z is the transmission line impedance. The noise spectrum at $\pm 2g$ drives bit flips within the dual-rail subspace with a combined T_1 limit of $T_1^{\text{DR}} = 2/S(2g)$ [55].

This noise affects both qubits according to their slopes $\partial \omega_T / \partial \Phi_e$ and their mutual inductances which are both $M \approx 1.28 \text{ pH}$. At the main operating point, Q1's slope vanishes, leading to a contribution just from Q2 which has slope $2\pi \times 7.8 \text{ GHz}/\Phi_0$. Using $T = 4 \text{ K}$ and $Z = 50\Omega$, this predicts a limit of $T_1^{\text{DR}} \approx 1 \text{ ms}$, which is roughly consistent with the measured T_1 . Further study is warranted to evaluate more directly whether this is the dominant error mechanism by altering the flux line configuration.

2. Photon fluctuations in readout resonators

Thermal fluctuations of photons in readout resonators have been known to limit qubit coherence in superconducting architectures. While we do not have a precise estimate for the thermal equilibrium population in the resonators, we can bound such population based on how it would limit the coherence of individual transmons. In particular, the transmon dephasing induced by thermal fluctuations in a coupled resonator is given by $\Gamma_\phi = 4\chi^2\eta\bar{n}/\kappa$, where χ is

the dispersive coupling between the transmon and the resonator, κ is the resonator linewidth, $\bar{n}$ is the thermal population, and $\eta = \kappa^2/(\kappa^2 + 4\chi^2)$ [57]. Using the readout parameters of Table I, we bound thermal population at $\bar{n} \lesssim 0.001$, which would cause a dephasing time of $T_\phi = 44 \mu\text{s}$ on Q1.

At this level of thermal population, we calculate how these fluctuations would affect dual-rail T_1^{DR} and T_2^{DR} . These fluctuations induce a Lorentzian noise spectrum on the qubits, $S(\omega) = 8\chi^2\kappa\bar{n}/(\omega^2 + \kappa^2)$ [57], whose contribution at the energy gap $S(2g)$ causes bit-flip errors in the dual-rail subspace. We calculate that this contribution for $\bar{n} = 0.001$ would impose $T_1^{\text{DR}} = 2/S(2g) \approx 10 \text{ ms}$ [55].

To estimate the induced dephasing on the dual-rail qubit, we numerically simulate the master equation for two coupled qubits, each dispersively coupled to a readout resonator with Hamiltonian,

$$H = g(\sigma_1^- \sigma_2^+ + \text{H.c.}) + \sum_{i=1,2} \chi_i \sigma_i^z a_i^\dagger a_i, \quad (\text{N3})$$

along with collapse operators $C_i^\downarrow = a_i \sqrt{\kappa_i(1 + \bar{n})}$ and $C_i^\uparrow = a_i^\dagger \sqrt{\kappa_i\bar{n}}$. In this model, we find that the induced dual-rail dephasing nearly saturates at the $2 \times T_1^{\text{DR}}$ limit imposed by the high-frequency part of the noise spectrum $S(2g)$. As such, we expect similarly that this dephasing is at a scale of $> 10 \text{ ms}$ and thus not limiting.

3. Interaction with TLSs

Parasitic two-level systems which are coupled to either of the two transmons comprising the dual-rail qubit can

affect its coherence. We consider a simple model with two transmons (described as bosonic modes a_1, a_2) and a TLS described by Pauli operators $\sigma_{\text{TLS}}^{x,y,z}$. While a TLS may dispersively shift its coupled transmon by a characteristic $\chi\sigma_z a^\dagger a$, it is insufficient to simply analyze this frequency shift in the context of dual-rail frequency sensitivity, as described by Eq. (M3). Instead, it is important to evaluate how the TLS couples to the dual-rail logical states.

We start with the following Hamiltonian in the rotating frame at the transmon frequency (excluding transmon anharmonicity for simplicity):

$$H_{\text{TLS}} = g(a_1^\dagger a_2 + \text{H.c.}) + \frac{\Delta}{2}\sigma_z + \lambda(a_1^\dagger \sigma_- + \text{H.c.}). \quad (\text{N4})$$

Rewriting in terms of the dual-rail eigenmodes $d_\pm = (1/\sqrt{2})(a_1 \pm a_2)$:

$$H = g(d_+^\dagger d_+ - d_-^\dagger d_-) + \frac{\Delta}{2}\sigma_z + \frac{\lambda}{\sqrt{2}}((d_+^\dagger + d_-^\dagger)\sigma_- + \text{H.c.}). \quad (\text{N5})$$

The TLS is now coupled equally to the two dual-rail eigenmodes at frequency $\pm g$, but possibly with a different detuning to the two modes.

Several effects may emerge from this model. Firstly, if the TLS has a short lifetime and is nearly resonant with one mode, it can cause one logical state to have a faster decay rate out of the dual-rail subspace. This is visible in Fig. 16 when the dual-rail qubit is parked at 4.96 GHz, while its upper mode is nearly resonant with a TLS which is observed directly on Q2 when parked at this upper hybrid mode frequency; specifically, we see that the upper mode has a short erasure lifetime.

Secondly, if not right on resonance with a dual-rail mode, the TLS becomes weakly hybridized with the modes, inducing dispersive shifts. Letting $\Delta_\pm = \Delta \pm g$ be the detuning to each mode, then the dispersive shifts are $\chi_\pm = \lambda^2/2\Delta_\pm$. The difference between these shifts, $\chi_{\text{DR}} = \chi_+ - \chi_-$, is what causes dual-rail dephasing, as this describes the effective dispersive coupling to the dual-rail qubit. Importantly, $|\chi_{\text{DR}}| = 2\lambda^2 g/(\Delta^2 - g^2)$ can be of the same order of magnitude as the dispersive coupling directly from the TLS onto a single isolated transmon, λ^2/Δ .

Finally, the hybridization between the dual-rail modes and the TLS can also cause the dual-rail qubit to inherit noise processes associated with the TLS. Slow switching can induce telegraph noise on the dual-rail qubit associated with toggling dispersive shifts, as observed experimentally and described in Fig. 13. (Notably, in this figure it is shown that the telegraph noise on the dual-rail qubit is of similar scale to that on the individual transmon, as compatible with the above analysis). Faster noise processes, including Markovian dephasing of the TLS, would include high-frequency noise components that can drive transitions

between logical states, limiting the dual-rail T_1^{DR} , as well as noise that contributes to direct dual-rail dephasing.

Further studies of how dual-rail qubits interact with TLSs are warranted, as this approach may offer complementary tools to characterize the properties of TLSs beyond what is available using single-transmon probes. For quantum computing applications, however, the dual-rail qubit primarily benefits by being able to tune its operating point to dodge TLSs.

APPENDIX O: SUPPLEMENTARY DATA ACROSS TUNABLE OPERATING POINTS

To supplement the core coherence metrics across the tunable operating range presented in Fig. 5, we show in Fig. 16 a more complete set of metrics across this range, including T_1, T_2^* , and T_2^{echo} for the individual transmons and the erasure lifetimes for the dual-rail qubit.

We note that the dip in dual-rail coherence which is observed at ~ 4.96 GHz is not matched with a dip in coherence in either Q1 or Q2 at those operating points. Instead, we attribute the dual-rail dip to a marked drop in Q2's T_1 at around 5.07 GHz. We hypothesize that these two dips are related because while the dual rail is operated at 4.96 GHz, its eigenmodes are present at $4.96 \text{ GHz} \pm 90 \text{ MHz}$, putting the upper mode near resonance with the TLS in Q2's spectrum. This is consistent with the notable reduction in the $|1_L\rangle$ -state erasure lifetime at this operating point.

APPENDIX P: CONTROL LINE AND CHIP RESOURCE ESTIMATION

While the dual-rail qubit offers significant advantages over standard transmons for quantum error correction, it comes at higher cost in terms of chip complexity and control lines. In this appendix, we enumerate such resource costs for two different possible dual-rail architectures and compare to two transmon-based architectures, with the summary presented in Table II. In all schemes, including dual rail and transmon based, we neglect resource costs for

TABLE II. Comparison of resource costs per qubit across transmon-based architectures. These resource costs are based on the discussion of each architecture in Appendix P.

Quantity per qubit	Dual-rail, ancilla qubit	Dual-rail, symmetric resonator	Tunable transmon	Fixed transmon
Readout resonators	3	1	1	1
XY lines	2	1	1	1
Z lines (total)	3	2	1	0
Z lines (fast flux)	1	1	1	0
Josephson junctions	6	4	2	1

tunable couplers, as similar coupler styles are possible in all architectures.

- (1) Dual rail with ancilla qubits. This scheme is the approach adopted in this paper, in which each dual-rail qubit is composed of two tunable transmons. An ancilla transmon (which could be fixed or tunable, but we assume tunable to maximize flexibility) which is coupled to this pair is used to detect erasure errors. All three transmons require a separate readout resonator. XY lines are required only for the ancilla and one of the dual-rail pair. Slow-flux tuning would be needed for all, but fast flux would be needed only on one of the dual-rail pair, both for modulation to drive single-qubit gates and also baseband to separate the pair for readout.
- (2) Dual rail with symmetric readout resonator. As described in Ref. [14], erasure detection can also be done using a readout resonator which is symmetrically coupled to the dual-rail pair. This same resonator can also be used to read out the final state of the dual-rail pair, either by shelving population outside of the dual-rail subspace (i.e., into $|11\rangle$) or by separating the transmons again using baseband flux. As a result, only one readout resonator is needed for the dual-rail qubit in this scheme, along with one XY line for either of the two transmons, slow flux for both, and fast flux for just one.
- (3) Tunable transmons. A tunable-transmon approach, as demonstrated in Ref. [1], involves one readout resonator for each transmon, as well as an XY line and fast-flux line, where in some cases the XY and Z signals can be combined.
- (4) Fixed transmons. Another standard fixed-transmon approach, as demonstrated in Ref. [58], similarly involves one readout resonator and XY line for each transmon, with no flux control.

[1] R. Acharya *et al.* (Google Quantum AI), *Suppressing quantum errors by scaling a surface code logical qubit*, *Nature (London)* **614**, 676 (2023).

[2] C. R. Clark, H. N. Tinkey, B. C. Sawyer, A. M. Meier, K. A. Burkhardt, C. M. Seck, C. M. Shappert, N. D. Guise, C. E. Volin, S. D. Fallek, H. T. Hayden, W. G. Rellergert, and K. R. Brown, *High-fidelity Bell-state preparation with $^{40}\text{Ca}^+$ optical qubits*, *Phys. Rev. Lett.* **127**, 130505 (2021).

[3] 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, J. A. Gerber, B. Neyenhuis, D. Hayes, and R. P. Stutz, *Realization of real-time fault-tolerant quantum error correction*, *Phys. Rev. X* **11**, 041058 (2021).

[4] D. M. Zajac, J. Stehlik, D. L. Underwood, T. Phung, J. Blair, S. Carnevale, D. Klaus, G. A. Keefe, A. Carniol, M. Kumph, M. Steffen, and O. E. Dial, *Spectator errors in tunable coupling architectures*, *arXiv:2108.11221*.

[5] S. J. Evered, D. Bluvstein, M. Kalinowski, S. Ebadi, T. Manovitz, H. Zhou, S. H. Li, A. A. Geim, T. T. Wang, N. Maskara, H. Levine, G. Semeghini, M. Greiner, V. Vuletić, and M. D. Lukin, *High-fidelity parallel entangling gates on a neutral-atom quantum computer*, *Nature (London)* **622**, 268 (2023).

[6] J. Guillaud and M. Mirrahimi, *Repetition cat qubits for fault-tolerant quantum computation*, *Phys. Rev. X* **9**, 041053 (2019).

[7] D. K. Tuckett, A. S. Darmawan, C. T. Chubb, S. Bravyi, S. D. Bartlett, and S. T. Flammia, *Tailoring surface codes for highly biased noise*, *Phys. Rev. X* **9**, 041031 (2019).

[8] A. S. Darmawan, B. J. Brown, A. L. Grimsmo, D. K. Tuckett, and S. Puri, *Practical quantum error correction with the XZZX code and Kerr-cat qubits*, *PRX Quantum* **2**, 030345 (2021).

[9] I. Cong, H. Levine, A. Keesling, D. Bluvstein, S.-T. Wang, and M. D. Lukin, *Hardware-efficient, fault-tolerant quantum computation with Rydberg atoms*, *Phys. Rev. X* **12**, 021049 (2022).

[10] C. Chamberland, K. Noh, P. Arrangoiz-Arriola, E. T. Campbell, C. T. Hann, J. Iverson, H. Putterman, T. C. Bohdanowicz, S. T. Flammia, A. Keller, G. Refael, J. Preskill, L. Jiang, A. H. Safavi-Naeini, O. Painter, and F. G. S. L. Brandão, *Building a fault-tolerant quantum computer using concatenated cat codes*, *PRX Quantum* **3**, 010329 (2022).

[11] 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*, *Nat. Phys.* **16**, 509 (2020).

[12] V. V. Sivak, A. Eickbusch, B. Royer, S. Singh, I. Tsioutsios, S. Ganjam, A. Miano, B. L. Brock, A. Z. Ding, L. Frunzio, S. M. Girvin, R. J. Schoelkopf, and M. H. Devoret, *Real-time quantum error correction beyond break-even*, *Nature (London)* **616**, 50 (2023).

[13] Y. Wu, S. Kolkowitz, S. Puri, and J. D. Thompson, *Erasure conversion for fault-tolerant quantum computing in alkaline earth Rydberg atom arrays*, *Nat. Commun.* **13**, 4657 (2022).

[14] A. Kubica, A. Haim, Y. Vaknin, H. Levine, F. Brandão, and A. Retzker, *Erasure qubits: Overcoming the T_1 limit in superconducting circuits*, *Phys. Rev. X* **13**, 041022 (2023).

[15] A. Kubica and A. Retzker, *Heralding of amplitude damping decay noise for quantum error correction*, U.S. Patent No. 11,748,652 (2021), <https://patents.google.com/patent/US11748652B1/en>.

[16] J. D. Teoh, P. Winkel, H. K. Babla, B. J. Chapman, J. Claes, S. J. De Graaf, J. W. O. Garmon, W. D. Kalfus, Y. Lu, A. Maiti, K. Sahay, N. Thakur, T. Tsunoda, S. H. Xue, L. Frunzio, S. M. Girvin, S. Puri, and R. J. Schoelkopf, *Dual-rail encoding with superconducting cavities*, *Proc. Natl. Acad. Sci. U.S.A.* **120**, e2221736120 (2023).

[17] M. Grassl, T. Beth, and T. Pellizzari, *Codes for the quantum erasure channel*, *Phys. Rev. A* **56**, 33 (1997).

[18] E. Knill, *Scalable quantum computation in the presence of large detected-error rates*, *arXiv:quant-ph/0312190*.

[19] K. Sahay, J. Jin, J. Claes, J. D. Thompson, and S. Puri, *High-threshold codes for neutral-atom qubits with biased erasure errors*, *Phys. Rev. X* **13**, 041013 (2023).

[20] W. C. Campbell, *Certified quantum gates*, *Phys. Rev. A* **102**, 022426 (2020).

- [21] M. Kang, W. C. Campbell, and K. R. Brown, *Quantum error correction with metastable states of trapped ions using erasure conversion*, [PRX Quantum](#) **4**, 020358 (2023).
- [22] S. Ma, G. Liu, P. Peng, B. Zhang, S. Jandura, J. Claes, A. P. Burgers, G. Pupillo, S. Puri, and J. D. Thompson, *High-fidelity gates and mid-circuit erasure conversion in an atomic qubit*, [Nature \(London\)](#) **622**, 279 (2023).
- [23] P. Scholl, A. L. Shaw, R. B.-S. Tsai, R. Finkelstein, J. Choi, and M. Endres, *Erasure conversion in a high-fidelity Rydberg quantum simulator*, [Nature \(London\)](#) **622**, 273 (2023).
- [24] I. L. Chuang and Y. Yamamoto, *Simple quantum computer*, [Phys. Rev. A](#) **52**, 3489 (1995).
- [25] R. Duan, M. Grassl, Z. Ji, and B. Zeng, *Multi-error-correcting amplitude damping codes*, in *Proceedings of the 2010 IEEE International Symposium on Information Theory, Austin, TX* (IEEE, New York, 2010), pp. 2672–2676.
- [26] E. Zakka-Bajjani, F. Nguyen, M. Lee, L. R. Vale, R. W. Simmonds, and J. Aumentado, *Quantum superposition of a single microwave photon in two different ‘colour’ states*, [Nat. Phys.](#) **7**, 599 (2011).
- [27] Y.-P. Shim and C. Tahan, *Semiconductor-inspired design principles for superconducting quantum computing*, [Nat. Commun.](#) **7**, 11059 (2016).
- [28] D. L. Campbell, Y.-P. Shim, B. Kannan, R. Winik, D. K. Kim, A. Melville, B. M. Niedzielski, J. L. Yoder, C. Tahan, S. Gustavsson, and W. D. Oliver, *Universal nonadiabatic control of small-gap superconducting qubits*, [Phys. Rev. X](#) **10**, 041051 (2020).
- [29] N. Timoney, I. Baumgart, M. Johanning, A. F. Varón, M. B. Plenio, A. Retzker, and C. Wunderlich, *Quantum gates and memory using microwave-dressed states*, [Nature \(London\)](#) **476**, 185 (2011).
- [30] K. C. Miao, J. P. Blanton, C. P. Anderson, A. Bourassa, A. L. Crook, G. Wolfowicz, H. Abe, T. Ohshima, and D. D. Awschalom, *Universal coherence protection in a solid-state spin qubit*, [Science](#) **369**, 1493 (2020).
- [31] Z. Huang, P. S. Mundada, A. Gyenis, D. I. Schuster, A. A. Houck, and J. Koch, *Engineering dynamical sweet spots to protect qubits from $1/f$ noise*, [Phys. Rev. Appl.](#) **15**, 034065 (2021).
- [32] T. Gullion, D. B. Baker, and M. S. Conradi, *New, compensated Carr-Purcell sequences*, [J. Magn. Reson.](#) **89**, 479 (1990).
- [33] D. C. McKay, C. J. Wood, S. Sheldon, J. M. Chow, and J. M. Gambetta, *Efficient Z gates for quantum computing*, [Phys. Rev. A](#) **96**, 022330 (2017).
- [34] W. P. Livingston, M. S. Blok, E. Flurin, J. Dressel, A. N. Jordan, and I. Siddiqi, *Experimental demonstration of continuous quantum error correction*, [Nat. Commun.](#) **13**, 2307 (2022).
- [35] 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**, 034040 (2018).
- [36] D. Ristè, C. C. Bultink, K. W. Lehnert, and L. DiCarlo, *Feedback control of a solid-state qubit using high-fidelity projective measurement*, [Phys. Rev. Lett.](#) **109**, 240502 (2012).
- [37] M. McEwen, D. Kafri, Z. Chen, J. Atalaya, K. J. Satzinger, C. Quintana, P. V. Klimov, D. Sank, C. Gidney, A. G. Fowler *et al.*, *Removing leakage-induced correlated errors in superconducting quantum error correction*, [Nat. Commun.](#) **12**, 1761 (2021).
- [38] J. Marques, H. Ali, B. Varbanov, M. Finkel, H. Veen, S. Van Der Meer, S. Valles-Sanclemente, N. Muthusubramanian, M. Beekman, N. Haider, B. Terhal, and L. DiCarlo, *All-microwave leakage reduction units for quantum error correction with superconducting transmon qubits*, [Phys. Rev. Lett.](#) **130**, 250602 (2023).
- [39] K. Khodjasteh and L. Viola, *Dynamically error-corrected gates for universal quantum computation*, [Phys. Rev. Lett.](#) **102**, 080501 (2009).
- [40] S. McArdle, X. Yuan, and S. Benjamin, *Error-mitigated digital quantum simulation*, [Phys. Rev. Lett.](#) **122**, 180501 (2019).
- [41] T.-Y. Yang, Y.-X. Shen, Z.-K. Cao, and X.-B. Wang, *Post-selection in noisy Gaussian boson sampling: Part is better than whole*, [Quantum Sci. Technol.](#) **8**, 045020 (2023).
- [42] D. R. Simon, *On the power of quantum computation*, [SIAM J. Comput.](#) **26**, 1474 (1997).
- [43] L. Botelho, A. Glos, A. Kundu, J. A. Miszczak, O. Salehi, and Z. Zimborás, *Error mitigation for variational quantum algorithms through mid-circuit measurements*, [Phys. Rev. A](#) **105**, 022441 (2022).
- [44] K. S. Chou, T. Shemma, H. McCarrick, T.-C. Chien, J. D. Teoh, P. Winkel, A. Anderson, J. Chen, J. Curtis, S. J. de Graaf *et al.*, *Demonstrating a superconducting dual-rail cavity qubit with erasure-detected logical measurements*, [arXiv:2307.03169](#).
- [45] S. Krinner, N. Lacroix, A. Remm, A. Di Paolo, E. Genois, C. Leroux, C. Hellings, S. Lazar, F. Swiadek, J. Herrmann, G. J. Norris, C. K. Andersen, M. Müller, A. Blais, C. Eichler, and A. Wallraff, *Realizing repeated quantum error correction in a distance-three surface code*, [Nature \(London\)](#) **605**, 669 (2022).
- [46] C. Macklin, K. O’Brien, D. Hover, M. E. Schwartz, V. Bolkhovskiy, X. Zhang, W. D. Oliver, and I. Siddiqi, *A near-quantum-limited Josephson traveling-wave parametric amplifier*, [Science](#) **350**, 307 (2015).
- [47] C. A. Ryan, B. R. Johnson, J. M. Gambetta, J. M. Chow, M. P. da Silva, O. E. Dial, and T. A. Ohki, *Tomography via correlation of noisy measurement records*, [Phys. Rev. A](#) **91**, 022118 (2015).
- [48] D. Sank, Z. Chen, M. Khezri, J. Kelly, R. Barends, 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*, [Phys. Rev. Lett.](#) **117**, 190503 (2016).
- [49] M. Khezri, A. Opremcak, Z. Chen, K. C. Miao, M. McEwen, A. Bengtsson, T. White, O. Naaman, D. Sank, A. N. Korotkov, Y. Chen, and V. Smelyanskiy, *Measurement-induced state transitions in a superconducting qubit: Within the rotating-wave approximation*, [Phys. Rev. Appl.](#) **20**, 054008 (2023).
- [50] K. Rudinger, S. Kimmel, D. Lobser, and P. Maunz, *Experimental demonstration of a cheap and accurate phase estimation*, [Phys. Rev. Lett.](#) **118**, 190502 (2017).

- [51] E. Lucero, J. Kelly, R. C. Bialczak, M. Lenander, M. Mariantoni, M. Neeley, A. D. O’Connell, D. Sank, H. Wang, M. Weides, J. Wenner, T. Yamamoto, A. N. Cleland, and J. M. Martinis, *Reduced phase error through optimized control of a superconducting qubit*, [*Phys. Rev. A* **82**, 042339 \(2010\)](#).
- [52] E. Knill, R. Laflamme, R. Martinez, and C.-H. Tseng, *An algorithmic benchmark for quantum information processing*, [*Nature \(London\)* **404**, 368 \(2000\)](#).
- [53] D. K. Weiss, H. Zhang, C. Ding, Y. Ma, D. I. Schuster, and J. Koch, *Fast high-fidelity gates for galvanically-coupled fluxonium qubits using strong flux modulation*, [*PRX Quantum* **3**, 040336 \(2022\)](#).
- [54] J. Bylander, S. Gustavsson, F. Yan, F. Yoshihara, K. Harrabi, G. Fitch, D. G. Cory, Y. Nakamura, J.-S. Tsai, and W. D. Oliver, *Noise spectroscopy through dynamical decoupling with a superconducting flux qubit*, [*Nat. Phys.* **7**, 565 \(2011\)](#).
- [55] The noise spectrum $S(\omega)$ drives bit flips within the dual-rail subspace with rates given by Fermi’s golden rule as $\Gamma_{0 \rightarrow 1} = \Gamma_{1 \rightarrow 0} = \frac{1}{4} S(2g)$, where the $1/4$ accounts for the matrix element between logical states [56]. Adding both bit-flip rates leads to a combined T_1 limit of $T_1^{\text{DR}} = 2/S(2g)$.
- [56] A. A. Clerk, M. H. Devoret, S. M. Girvin, F. Marquardt, and R. J. Schoelkopf, *Introduction to quantum noise, measurement, and amplification*, [*Rev. Mod. Phys.* **82**, 1155 \(2010\)](#).
- [57] P. Krantz, M. Kjaergaard, F. Yan, T. P. Orlando, S. Gustavsson, and W. D. Oliver, *A quantum engineer’s guide to superconducting qubits*, [*Appl. Phys. Rev.* **6**, 021318 \(2019\)](#).
- [58] P. Jurcevic, A. Javadi-Abhari, L. S. Bishop, I. Lauer, D. F. Bogorin, M. Brink, L. Capelluto, O. Günlük, T. Itoko, N. Kanazawa *et al.*, *Demonstration of quantum volume 64 on a superconducting quantum computing system*, [*Quantum Sci. Technol.* **6**, 025020 \(2021\)](#).