

Experimental Quantum Error Correction below the Surface Code Threshold via All-Microwave Leakage Suppression

Tan He^{1,2,3,*}, Weiping Lin^{1,2,3,*}, Rui Wang^{1,2,3,*}, Yuan Li^{1,2,3,*}, Jiahao Bei,² Jianbin Cai,^{1,2,3} Sirui Cao,^{1,2,3} Danning Chen,⁴ Kefu Chen^{1,2,3}, Xiawei Chen², Zhe Chen,⁴ Zhiyuan Chen^{1,2,3}, Zihua Chen,^{1,2,3} Wenhao Chu⁴, Hui Deng^{1,2,3}, Xun Ding,³ Zhuzhengqi Ding,² Bo Fan,² Daojin Fan,^{1,2,3} Yuanhao Fu,^{1,2,3} Dongxin Gao,^{1,2,3} Ming Gong,^{1,2,3} Jiacheng Gui³, Cheng Guo^{1,2,3}, Shaojun Guo^{1,2,3}, Lianchen Han,^{1,2,3} Linyin Hong,⁴ Yisen Hu,^{1,2,3} He-Liang Huang,⁵ Yong-Heng Huo^{1,2,3}, Chenyan Jiang,³ Lei Jiang,^{1,2,3} Tao Jiang^{1,2,3}, Zuokai Jiang,² Honghong Jin,² Dayu Li,^{1,2,3} Dongdong Li,⁴ Jiaqi Li,² Jinjin Li³, Junyan Li,² Junyun Li^{1,2,3}, Na Li,^{1,2,3} Shaowei Li,^{1,2,3} Yuhuai Li,^{1,2,3} Futian Liang^{1,2,3}, Nanxing Liao,² Jin Lin,^{1,2,3} Ke Liu^{1,2,3}, Maliang Liu,⁶ Yancheng Liu,^{1,2,3} Haoxin Lou,² Yuwei Ma^{1,2,3}, Kailiang Nan,³ Meijuan Nie,² Le Niu², Wenyi Peng³, Haoran Qian^{1,2,3}, Hao Rong,^{1,2,3} Tao Rong,^{1,2,3} Huiyan Shen,⁴ Qiong Shen,² Hong Su,^{1,2,3} Feifan Su,^{1,2,3} Chenyin Sun,^{1,2,3} Liangchao Sun,⁴ Tianzuo Sun,^{1,2,3} Yingxiu Sun,⁴ Yimeng Tan,² Jun Tan³, Wenbing Tu,⁴ Jiafei Wang,⁴ Biao Wang,⁴ Chang Wang,⁴ Chen Wang,^{1,2,3} Chu Wang,^{1,2,3} Jian Wang,³ Shengtao Wang,³ Xinzhe Wang³, Zuolin Wei^{1,2,3}, Dachao Wu,^{1,2,3} Gang Wu,^{1,2,3} Yulin Wu^{1,2,3}, Yu Xu^{1,2,3}, Chun Xue,⁴ Kai Yan,^{1,2,3} Xin Yan,⁴ Weifeng Yang,⁴ Xinpeng Yang^{1,2,3}, Yang Yang,² Yangsen Ye,^{1,2,3} Zhenping Ye,^{1,2,3} Zhengzhong Yi,^{1,2,3} Chong Ying,^{1,2,3} Jiale Yu,^{1,2,3} Qijing Yu^{1,2,3}, Xiangdong Zeng³, Chen Zha,^{1,2,3} Shaoyu Zhan,^{1,2,3} Haibin Zhang,³ He Zhang,³ Kaili Zhang,² Wen Zhang,² Yiming Zhang,^{1,2,3} Yongzhuo Zhang,³ Ziyang Zhang,⁴ Guming Zhao,^{1,2,3} Xintao Zhao,² Youwei Zhao,^{1,2,3} Zhong Zhao,⁴ Luyuan Zheng,² Fei Zhou⁷, Liang Zhou,⁴ Na Zhou,² Naibin Zhou^{1,2,3}, Chengjun Zhu,³ Qingling Zhu,^{1,2,3} Guihong Zou³, Haonan Zou,² Qiang Zhang,^{1,2,3,7} Chao-Yang Lu^{1,2,3}, Cheng-Zhi Peng,^{1,2,3} Fusheng Chen,^{1,2,3} XiaoBo Zhu^{1,2,3,7} and Jian-Wei Pan^{1,2,3}

¹*Hefei National Research Center for Physical Sciences at the Microscale and School of Physical Sciences, University of Science and Technology of China, Hefei 230026, China*

²*Shanghai Research Center for Quantum Science and CAS Center for Excellence in Quantum Information and Quantum Physics, University of Science and Technology of China, Shanghai 201315, China*

³*Hefei National Laboratory, University of Science and Technology of China, Hefei 230088, China*

⁴*QuantumCTek Co., Ltd., Hefei 230026, China*

⁵*Henan Key Laboratory of Quantum Information and Cryptography, Zhengzhou, Henan 450000, China*

⁶*School of Microelectronics, Xidian University, Xi'an, China*

⁷*Jinan Institute of Quantum Technology and Hefei National Laboratory Jinan Branch, Jinan 250101, China*



(Received 9 October 2025; revised 7 November 2025; accepted 25 November 2025; published 22 December 2025)

Quantum error correction (QEC) enables practical quantum computing by encoding logical qubits in many physical qubits, which can exponentially suppress the logical error rate with increasing code size provided that the physical error rate is below a critical threshold. However, the leakage of quantum information from the computational subspace presents a critical challenge to the development of scalable QEC, which creates long-lived, correlated errors that spread across space and time. Here, we demonstrate a quantum memory operating below the threshold by implementing an all-microwave leakage suppression architecture on a distance-7 surface code. We achieve a logical error suppression factor of $\Lambda = 1.40(6)$, definitively reversing the above-threshold scaling ($\Lambda < 1$) caused by unmitigated leakage. This scheme integrates a hardware-efficient leakage reduction unit for data qubits with a fast, unconditional reset for ancilla qubits, suppressing the average leakage population after 40 cycles by a factor of 72 to $6.4(5) \times 10^{-4}$. Our results demonstrate the viability of all-microwave control architectures for suppressing critical errors at scale, paving the way for more advanced quantum error correction implementations.

DOI: 10.1103/rqkg-dw31

Quantum error correction (QEC) provides the essential pathway toward fault-tolerant quantum computation [1–3]. The central goal of QEC is to operate below a critical

threshold, a regime where the logical error rate can be exponentially suppressed with increasing code distance. Achieving this on a physical device presents a formidable challenge: the processor's scale must be increased while simultaneously suppressing the error rates of all component

*These authors contributed equally to this work.

operations to exceedingly low levels. The surface code [4], with its high threshold and nearest-neighbor interactions, is the leading candidate for experimental QEC. Early proof-of-principle surface-code experiments on superconducting processors—including the distance-3 demonstrations [5,6] and subsequent distance-5 implementations [7]—validated the architecture but did not yet achieve below-threshold performance, being limited by the combined impact of gate and measurement infidelities, residual crosstalk, and other noise mechanisms. As the system grows in size and circuit depth, leakage—the population outside the computational subspace—becomes a particularly damaging correlated error source that violates the assumption of independence of physical errors in QEC, thereby degrading the exponential suppression of logical errors [8].

Physically mitigating leakage requires active removal of the population from noncomputational states. Two general strategies have been developed. One strategy leverages baseband flux pulses to dynamically shift the qubit frequency into resonance with a lossy readout resonator, thereby inducing a SWAP process that dissipates the leaked population [9–11]. We refer to this as a dc-based approach. Recent work showed that, as the surface code distance increases, this approach substantially lowers the logical error rate by effectively suppressing leakage [12]. However, this technique is generally limited to specific device architectures, most notably the requirement that resonator frequencies must lie below those of the qubits [9,10]. This requirement not only makes the system prone to readout-induced transitions but also significantly restricts processor architecture design [13]. An alternative paradigm, which we adopt in this Letter, overcomes these constraints. It is an all-microwave approach analogous to sideband cooling [14], which utilizes microwave drives to activate specific sideband transitions that efficiently return leaked population to the computational subspace [15,16]. Such a strategy offers hardware efficiency and removes strict constraints on device frequencies. Moreover, its inherent frequency-multiplexing capability is expected to reduce wiring complexity in cryogenic environments [17,18]. Although preliminary demonstrations on small-scale systems have shown promising results [19–22], the performance of this approach remains unvalidated in full-scale QEC architectures, where frequency crowding and microwave crosstalk emerge as key challenges.

In this Letter, we bridge this critical gap by demonstrating a scalable, all-microwave leakage suppression architecture in complete surface code implementations [4,23]. This integrated scheme combines leakage reduction units (LRUs) on data qubits with unconditional RST on ancilla qubits. Crucially, this integrated design is executed without incurring additional overhead in either time or space, as it requires neither dedicated operational stages nor dedicated leakage-removal circuit elements. By implementing an efficient calibration protocol with optimized frequency

allocation, we successfully suppress *in situ* leakage by a factor of 72 at the distance-7 surface code. This suppression is the key factor enabling our system to surpass the fault-tolerance threshold, demonstrated by a logical error suppression factor of $\Lambda = 1.40(6)$. Our Letter, together with recent below-threshold demonstrations using a dc-based leakage suppression scheme [12,24], collectively establishes a unifying principle for large-scale QEC: suppressing leakage is the critical requirement for below-threshold performance. These combined results provide a robust blueprint for realizing more advanced fault-tolerant quantum computation.

Our experiments are performed on the *Zuchongzhi 3.2* processor, featuring frequency-tunable transmon qubits arranged in a square lattice [25,26]. The architecture employs tunable couplers for nearest-neighbor connectivity, specifically optimized for surface code operations. We utilize 97 selected qubits capable of implementing both a complete distance-7 surface code and multiple independent distance-5 and distance-3 surface codes [Fig. 1(a)]. The processor is characterized by an average single-qubit gate error of 0.089%, two-qubit gate error of 0.543%, and readout error of 0.95% (see the Supplemental Material [27]).

Surface code operation relies on repeated syndrome extraction cycles, where each cycle's operations, including single-qubit gates [41], CZ gates [42,43], and measurements [13,44], can induce leakage errors. These errors accumulate both spatially (through qubit interactions) and temporally (across cycles), progressively degrading logical performance [10]. A key innovation of our scheme is the integration of LRU and RST directly into the surface code cycle circuit, as illustrated by the timing diagram in Fig. 1(d). This concurrent design is highly efficient. It avoids introducing any temporal overhead, unlike schemes requiring an additional data qubit leakage removal operation [10], and critically, requires no spatial overhead from dedicated leakage-removal qubits [12] while also utilizing the same resonators for both qubit readout and the dissipative leakage suppression operations. This enables a minimal QEC cycle, which is crucial for improving logical performance by both increasing the effective clock speed and limiting the accumulation of idle errors.

The all-microwave leakage suppression scheme is based on parametric flux modulation, similar to previous work in Refs. [21,45], which activates specific sideband transitions as depicted in the energy level diagram and operational schematic in Figs. 1(b) and 1(c). The physical mechanism relies on two complementary processes tailored for different qubit types. For data qubits, a single-step parametric drive ($|f0\rangle \xrightarrow{\omega_d} |e1\rangle \xrightarrow{\kappa_r} |e0\rangle$), designated as the *f*-LRU, selectively returns leakage population to the computational subspace while preserving encoded information. Ancilla qubits utilize a two-step process: first converting $|f0\rangle$ to $|e0\rangle$ through the same *f*-LRU transition, followed by an

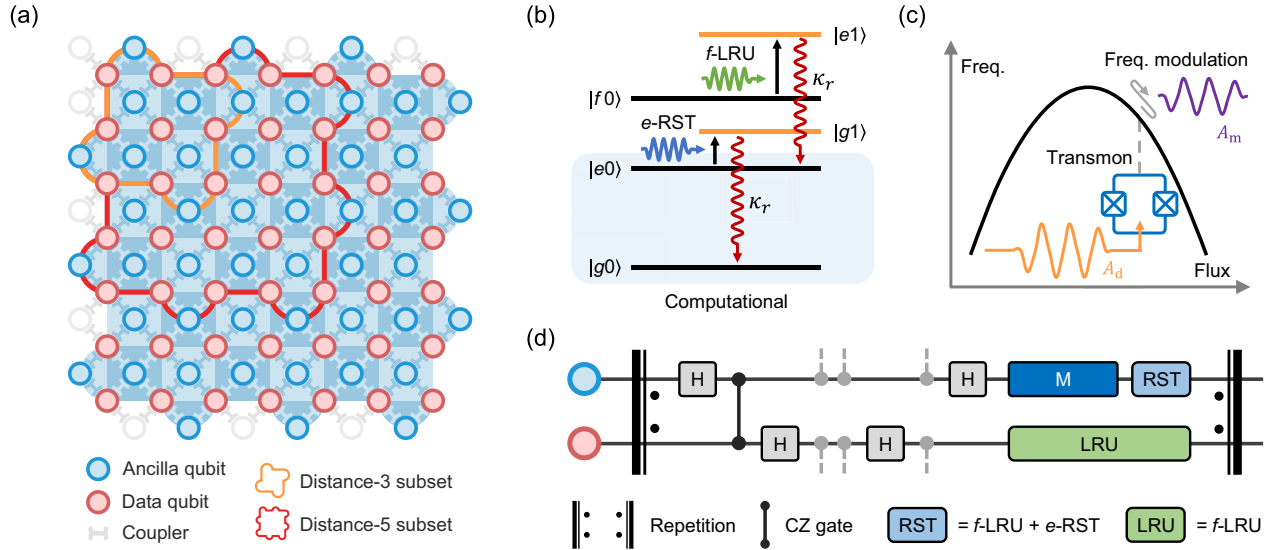


FIG. 1. Quantum processor architecture and all-microwave leakage reduction scheme. (a) Schematic of the quantum processor. We selected 97 qubits from a 107-qubit processor, enabling the implementation of a full distance-7 surface code. (b) Energy-level diagram of the leakage suppression mechanism. Microwave drives activate two key transitions: an f -LRU ($|f0\rangle \leftrightarrow |f1\rangle$) to remove leakage from the $|f\rangle$ state, and an e -RST ($|e0\rangle \leftrightarrow |g1\rangle$) to deplete the $|e\rangle$ state. Population is removed from the system via the resonator's fast decay rate, κ_r (typical ringdown time ≈ 15 ns). (c) Principle of parametric flux modulation. A microwave drive with amplitude A_d modulates the flux of the transmon qubit, resulting in a modulation of its frequency with amplitude A_m . (d) A single cycle of the surface code stabilizer measurement circuit. A representative data-ancilla pair is shown. The LRU on the data qubit is applied concurrently with the measurement (M) and RST on the ancilla qubit, while the data qubit is simultaneously protected by a dynamical decoupling sequence, to suppress leakage without increasing the cycle time. See the Supplemental Material [27] for detailed pulse sequences.

additional e -RST step ($|e0\rangle \xrightarrow{\omega_d} |g1\rangle \xrightarrow{\kappa_r} |g0\rangle$) to achieve complete RST. In this notation, the first label (g, e, f) represents the transmon energy state (ground, excited, leakage), while the second digit (0, 1) corresponds to the resonator photon number. κ_r is the resonator energy-decay rate that irreversibly drains excitation into the environment. The drive ω_d is chosen to satisfy the specific dressed resonance.

Scaling this all-microwave approach to a full-scale QEC architecture requires a carefully designed calibration protocol that mitigates challenges inherent to large systems, namely microwave crosstalk and frequency crowding. To address this, we developed an efficient and scalable calibration scheme. First, we improve our frequency arrangement strategy that treats the LRU and RST drive frequencies as new constraints, optimizing them alongside the existing frequencies for single-qubit gates, two-qubit gates, and readout to avoid collisions (see the Supplemental Material [27]). Second, we re-engineered the RST protocol's temporal structure by decoupling the $|f0\rangle \rightarrow |e1\rangle \rightarrow |e0\rangle$ and $|e0\rangle \rightarrow |g1\rangle \rightarrow |g0\rangle$ transitions into distinct sequential phases. While mathematically equivalent to simultaneous driving [45], this approach reduces waveform complexity by enabling the independent tuning of control parameters, which is critical as system size scales.

Experimentally, we first calibrate the RST protocol for the ancilla qubits, which consists of a 44 ns f -LRU pulse followed by a 44 ns e -RST pulse (total duration 96 ns,

including 4 ns padding at the start and end). The calibration for each pulse is a two-stage process, detailed in Figs. 2(a) and 2(b). The procedure begins by preparing the qubit in its corresponding initial state ($|f\rangle$ or $|e\rangle$) and then applying the test drive pulse. First, we perform a coarse sweep of the drive amplitude and frequency to map the amplitude-dependent frequency shift caused by the ac Stark effect, fitting the optimal frequency to a quadratic function [45]. Second, using this fitted relationship, we perform a fine sweep of the drive amplitude to find the value that minimizes the residual population in the initial state ($|f\rangle$ or $|e\rangle$). The first minimum of the resulting population oscillation is chosen as the optimal operating point. This RST protocol is also implemented on data qubits to accelerate initialization, reducing the shot-to-shot repetition time to 100 μ s.

Next, we calibrate the LRU for the data qubits. The LRU is a single 496 ns f -LRU pulse, framed by 24 ns padding for a total duration of 544 ns to match the ancilla's combined measurement and RST time [Fig. 1(d)]. While the calibration protocol begins similarly to the ancilla's, the significantly longer pulse duration alters the underlying dynamics. The system enters a regime where the residual $|f\rangle$ state population decays monotonically with drive amplitude instead of oscillating, thus exhibiting no clear minimum [Fig. 2(a)]. We therefore adopt a new criterion: setting the drive parameters to achieve a target leakage

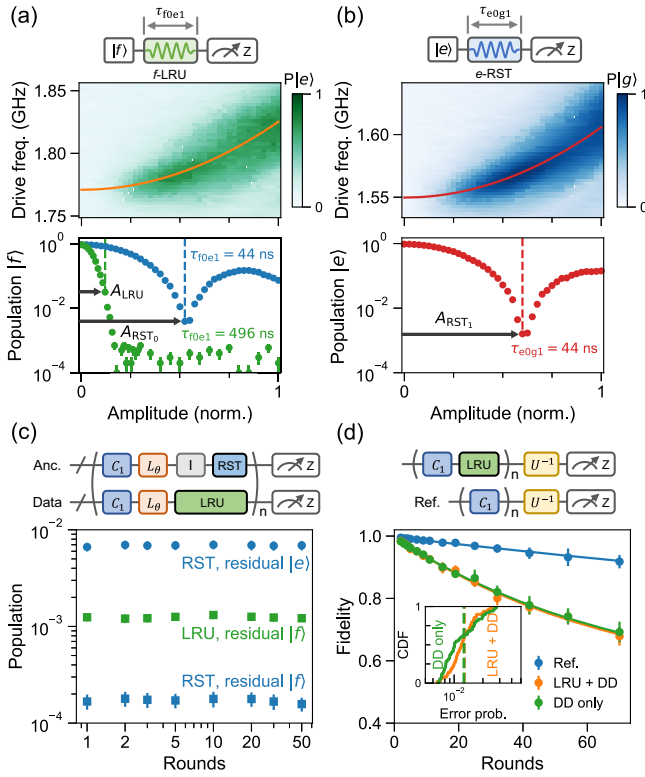


FIG. 2. Calibration and benchmarking for RST and LRU. (a) Schematic of f -LRU calibration. Top: pulse sequence used for f -LRU calibration. Middle: a coarse frequency-amplitude sweep to characterize the ac Stark shift. For each amplitude, we identify the frequency that maximizes the target probability and perform a quadratic fit (orange lines). Bottom: residual $|f\rangle$ population versus drive amplitude. A fine amplitude sweep along the fitted quadratic curve to find the optimal drive amplitude, using pulses of two fixed durations: 44 ns for ancilla qubits and 496 ns for data qubits. (b) Schematic of e -RST calibration. The procedure mirrors (a) with the qubit initialized in $|e\rangle$; the e -RST is calibrated by tracking the depletion of the $|f\rangle$ population. (c) Performance of the parallel RST and LRU protocols, benchmarked simultaneously across all 97 qubits using a randomized sequence with repeated 5% leakage injection. The plotted curves show the average residual populations for these two respective groups. The residual populations are slightly overestimated due to the inclusion of readout separation error. (d) Characterization of computational error from the LRU via the interleaved RB protocol. Each point is computed from 25 random sequences and fit with an exponential decay. The reference curve characterizes the intrinsic error rate of the single-qubit Clifford gates. Inset: error distributions for LRU + DD and DD-only protocols across 49 data qubits. The result shows near-identical error probability distributions, confirming that the LRU adds negligible computational error.

reduction. A target of 95% is sufficient, as the exponential nature of leakage suppression renders the benefit of higher reductions marginal [46].

Following the calibration of optimal drive parameters, we evaluated the performance of the RST and LRU protocols by executing a parallel randomized benchmarking (RB)-like

sequence simultaneously across all 97 qubits of the distance-7 code, depicted in Fig. 2(c). Each cycle of this sequence consists of a random single-Clifford gate, a subsequent pulse (L_θ) that injects approximately 5% leakage into the $|f\rangle$ state, and the final suppression operation: the RST protocol was applied to all 48 ancilla qubits, while the LRU protocol was applied concurrently to all 49 data qubits. For the ancilla qubits, the RST protocol rapidly suppresses the residual populations in both the $|e\rangle$ and $|f\rangle$ states, which saturate below 7×10^{-3} and 2×10^{-4} , respectively. Concurrently, the LRU protocol on the data qubits reduces the $|f\rangle$ state population to below 2×10^{-3} . The rapid saturation at these low population levels demonstrates the efficacy of our protocols in providing leakage suppression across the entire device.

A crucial requirement for any leakage suppression technique is that it introduces negligible error into the computational subspace [8]. To verify this, we performed simultaneous interleaved RB to precisely quantify the error added by the LRU operation [47], as shown in Fig. 2(d). Since the LRU and dynamical decoupling (DD) are applied concurrently in the surface code, we isolate the LRU's contribution to the error by comparing the fidelity of an LRU + DD sequence to a DD-only baseline. The measured error rates of 1.28(2)% (LRU + DD) and 1.27(2)% (DD only) are statistically identical. This negligible difference demonstrates that the LRU introduces negligible error into the computational subspace, confirming its compatibility with high-fidelity quantum error correction.

Having established the efficacy of the suppression protocols at the component level, we next evaluated their system-level impact during the execution of a full distance-7 surface code. As shown in Fig. 3(a), we tracked the average leakage population over 40 QEC cycles for four configurations: full mitigation (“Both”), partial mitigation (“LRU only” and “RST only”), and no mitigation (“Neither”). Without any intervention, the leakage exhibits uncontrolled accumulation, reaching $4.64(4) \times 10^{-2}$ by cycle 40. The partial schemes proved insufficient, as applying only RST or only LRU selectively suppressed leakage on just the ancilla or data qubits, respectively. In stark contrast, the combined strategy suppresses leakage to a steady-state population of $6.4(5) \times 10^{-4}$ at cycle 40—a 72-fold improvement over the unmitigated baseline. This temporal dynamic is further explained by the spatial maps in Fig. 3(b), which reveal that without mitigation, leakage not only grows in time but also spreads across the device, creating localized hotspots that are fatal to QEC performance [10]. Our combined RST + LRU strategy effectively suppresses both the temporal accumulation and spatial diffusion of leakage.

Furthermore, we evaluate our architecture's impact on logical performance through quantum memory experiments [5–7, 12, 24]. We first prepare the logical state in the X or Z basis, run a variable number of QEC cycles, and finally

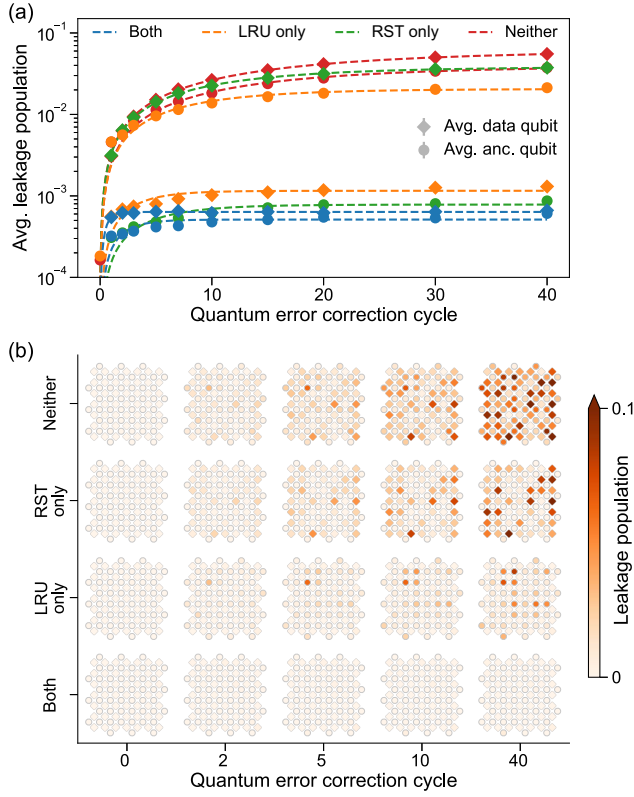


FIG. 3. Leakage suppression during distance-7 surface code execution. (a) Average leakage population versus QEC cycle for four mitigation strategies: full suppression with both RST and LRU (Both), partial suppression (LRU only, RST only), and no suppression (Neither). Data qubits and ancilla qubits are represented by diamonds and circles, respectively. (b) Spatial maps of leakage population across the device. Full leakage suppression prevents the temporal and spatial growth of leakage, eliminating the high-leakage hotspots that otherwise emerge after multiple cycles.

measure the logical qubit in the initialized logical basis. The resulting syndromes are processed by decoders using priors derived from a comprehensive, circuit-level noise model that explicitly includes all operation-specific error channels. In the experiments, we explored several algorithms, including a matching decoder [48] and k -iteration correlated-matching decoder and belief-matching decoder [49], which are detailed in the Supplemental Material [27]. While the correlated-matching decoder offers a better balance of speed and accuracy, we used the belief-matching decoder for the results shown in Fig. 4 to achieve the highest logical fidelity.

As shown for the distance-7 code in Fig. 4(a), the logical fidelity is dramatically improved with full leakage mitigation (Both). Crucially, its decay follows a single-exponential trend, indicating a stable logical error rate per cycle, ϵ_L . In contrast, configurations lacking full mitigation exhibit a faster, nonexponential decay. This is a direct signature of an error rate that increases over time as leakage accumulates, rapidly degrading the logical state.

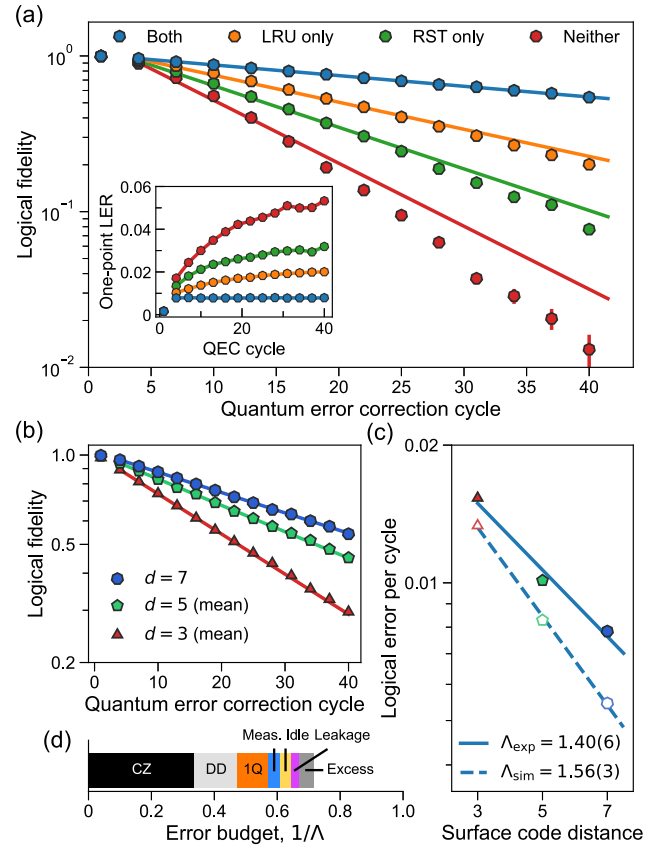


FIG. 4. Leakage mitigation improves surface code logical performance. (a) Logical fidelity versus QEC cycles for a distance-7 surface code under four different mitigation schemes. Each data point represents 5×10^4 repetitions and is averaged over X and Z bases. The solid lines represent exponential fits to the experimental data. Notably, only the “Both” condition (blue) follows a single-exponential trend, while the other schemes show a faster, nonexponential decay. The inset shows the corresponding one-point logical error rate per cycle (LER) versus cycles for each scheme. (b) Logical fidelity versus cycles for surface codes with distance 3, 5, and 7, all with full leakage suppression applied. (c) The extracted logical error per cycle as a function of surface code distance. The solid line is a fit to the experiment data using the model $\epsilon_L \propto \Lambda^{-(d+1)/2}$, yielding a suppression factor of $\Lambda_{\text{exp}} = 1.40(6)$, and the dashed line is a fit to the simulated data, yielding $\Lambda_{\text{sim}} = 1.56(3)$. (d) The estimated error budget, showing the contributions of different error sources based on component-level characterization and simulation. See the Supplemental Material [27] for details.

Furthermore, we extract ϵ_L with a single number of cycles [12]. As plotted in the inset, we find that ϵ_L remains stable only under full leakage mitigation. For all other configurations, ϵ_L increases significantly with the number of cycles.

We next characterize the scaling of the logical error with code distance under full leakage mitigation. To ensure a fair comparison, the performance of nine distance-3 and four distance-5 subgrids is averaged and compared with that of the full distance-7 code, with all error rates averaged over

the logical Pauli X_L and Z_L bases. Figure 4(b) shows the logical error performance of the different code distance. By fitting the logical error probability to an exponential decay, we observe a monotonic decrease in ε_L with increasing code size, as plotted in Fig. 4(c). We measure $\varepsilon_{L,d=3} = 1.532(2)\%$, $\varepsilon_{L,d=5} = 1.012(2)\%$, and $\varepsilon_{L,d=7} = 0.783(3)\%$. This trend is the experimental signature of operating below threshold for our specific implementation. Fitting these data to the model $\varepsilon_L \propto \Lambda^{-(d+1)/2}$ yields a logical error suppression factor of $\Lambda = 1.40(6)$, which represents the suppression of the logical error ε_L for each step of increasing distance from d to $d+2$. This experimental value is slightly lower than the $\Lambda_{\text{sim}} = 1.56(3)$ obtained from our simulation. We attribute this modest discrepancy to unmodeled error sources in the simulation, such as residual crosstalk and other nonlocal noise mechanisms. Conversely, in the absence of full leakage suppression, we find this scaling is inverted ($\Lambda < 1$), leading to error amplification where larger codes perform worse (see the Supplemental Material [27]). Finally, the error budget analysis shown in Fig. 4(d) provides critical insight into the system's error hierarchy [7]. With our effective mitigation protocols, leakage is no longer a dominant error contributor. The budget is now limited primarily by the infidelity of our two-qubit CZ gates.

In summary, we have demonstrated a surface code operating below the error threshold for distances up to $d = 7$, achieving a logical error suppression factor of $\Lambda = 1.40(6)$. This performance was enabled by a scalable, hardware-efficient all-microwave leakage suppression architecture, integrating both LRU and RST operations, which successfully suppressed leakage by a factor of 72 at the distance-7 surface code. The key to implementing this architecture at scale was an efficient calibration protocol with optimized frequency allocation, designed to mitigate the frequency crowding and microwave crosstalk inherent to a large-scale parallel operation.

Though demonstrated with the surface code, our architecture is broadly compatible with other QEC frameworks [50,51], as it operates independently of the specific stabilizer measurements. Looking ahead, integrating this architecture with all-microwave two-qubit gates [42] and frequency multiplexing promises a scalable, fully microwave-driven quantum computer architecture with significantly reduced hardware overhead.

Acknowledgments—This research was supported by the Innovation Program for Quantum Science and Technology-National Science and Technology Major Project (Grant No. 2021ZD0300200), National Natural Science Foundation of China (Grant No. 92476203 and No. 92476001), Anhui Initiative in Quantum Information Technologies, the Special funds from Jinan Science and Technology Bureau and Jinan High Tech Zone Management Committee, Shanghai Municipal Science and Technology

Major Project (Grant No. 2019SHZDZX01), and the XPLOER PRIZE.

X. Z., C.-Z. P., and J.-W. P. conceived the research. T. H. and W. L. conducted the experiment under the supervision of F. C. The gate calibration was performed by Z. C. and T. J. D. G., R. W., S. G., and K. C. performed the readout calibration. Y. L. and S. C. fabricated the device. Z. Y., X. Y., and Y. Z. developed decoders. X. Z., L. N., and X. C. provided software tools. Y. M., Y. Y., S. G., and H. Q. built the test systems. J. L., C. G., F. L., and C.-Z. P. developed the electronics. F. C., H. H., T. H., and W. L. wrote the manuscript, with inputs from all coauthors. All authors contributed to the discussion of the results and the development of this work. X. Z., C.-Z. P., and J.-W. P. supervised the whole project.

Data availability—The data that support the findings of this article are not publicly available upon publication because it is not technically feasible and/or the cost of preparing, depositing, and hosting the data would be prohibitive within the terms of this research project. The data are available from the authors upon reasonable request.

-
- [1] P. W. Shor, *Phys. Rev. A* **52**, R2493 (1995).
 - [2] E. Dennis, A. Kitaev, A. Landahl, and J. Preskill, *J. Math. Phys.* (N.Y.) **43**, 4452 (2002).
 - [3] B. M. Terhal, *Rev. Mod. Phys.* **87**, 307 (2015).
 - [4] A. G. Fowler, M. Mariantoni, J. M. Martinis, and A. N. Cleland, *Phys. Rev. A* **86**, 032324 (2012).
 - [5] Y. Zhao, Y. Ye, H.-L. Huang, Y. Zhang, D. Wu, H. Guan, Q. Zhu, Z. Wei, T. He, S. Cao *et al.*, *Phys. Rev. Lett.* **129**, 030501 (2022).
 - [6] S. Krinner, N. Lacroix, A. Remm, A. D. Paolo, E. Genois, C. Leroux, C. Hellings, S. Lazar, F. Swiadek, J. Herrmann *et al.*, *Nature (London)* **605**, 669 (2022).
 - [7] Google Quantum AI, *Nature (London)* **614**, 676 (2023).
 - [8] P. Aliferis and B. M. Terhal, *Quantum Inf. Comput.* **7**, 139 (2005).
 - [9] M. McEwen, D. Kafri, Z. Chen, J. Atalaya, K. Satzinger, C. Quintana, P. V. Klimov, D. Sank, C. Gidney, A. Fowler *et al.*, *Nat. Commun.* **12**, 1761 (2021).
 - [10] K. C. Miao, M. McEwen, J. Atalaya, D. Kafri, L. P. Pryadko, A. Bengtsson, A. Opremcak, K. J. Satzinger, Z. Chen, P. V. Klimov *et al.*, *Nat. Phys.* **19**, 1780 (2023).
 - [11] X. Yang, J. Chu, Z. Guo, W. Huang, Y. Liang, J. Liu, J. Qiu, X. Sun, Z. Tao, J. Zhang *et al.*, *Phys. Rev. Lett.* **133**, 170601 (2024).
 - [12] Google Quantum AI, *Nature (London)* **638**, 920 (2025).
 - [13] 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, *Phys. Rev. Appl.* **20**, 054008 (2023).
 - [14] C. Monroe, D. M. Meekhof, B. E. King, S. R. Jefferts, W. M. Itano, D. J. Wineland, and P. Gould, *Phys. Rev. Lett.* **75**, 4011 (1995).

- [15] A. Blais, J. Gambetta, A. Wallraff, D. I. Schuster, S. M. Girvin, M. H. Devoret, and R. J. Schoelkopf, *Phys. Rev. A* **75**, 032329 (2007).
- [16] F. Beaudoin, M. P. da Silva, Z. Dutton, and A. Blais, *Phys. Rev. A* **86**, 022305 (2012).
- [17] B. Patra, R. M. Incandela, J. P. G. van Dijk, H. A. R. Homulle, L. Song, M. Shahmohammadi, R. B. Staszewski, A. Vladimirescu, M. Babaie, F. Sebastiano, and E. Charbon, *IEEE J. Solid-State Circuits* **53**, 309 (2018).
- [18] R. Huang, X. Geng, X. Wu, G. Dai, L. Yang, J. Liu, and W. Chen, *Phys. Rev. Appl.* **18**, 064046 (2022).
- [19] P. Magnard, P. Kurpiers, B. Royer, T. Walter, J.-C. Besse, S. Gasparinetti, M. Pechal, J. Heinsoo, S. Storz, A. Blais, and A. Wallraff, *Phys. Rev. Lett.* **121**, 060502 (2018).
- [20] J. F. Marques, H. Ali, B. M. Varbanov, M. Finkel, H. M. Veen, S. L. M. van der Meer, S. Valles-Sanclemente, N. Muthusubramanian, M. Beekman, N. Haider *et al.*, *Phys. Rev. Lett.* **130**, 250602 (2023).
- [21] N. Lacroix, L. Hofele, A. Remm, O. Benhayoune-Khadraoui, A. McDonald, R. Shillito, S. Lazar, C. Hellings, F. M. C. Swiadek, D. Colao-Zanuz, A. Flasby, M. B. Panah, M. Kerschbaum, G. J. Norris, A. Blais, A. Wallraff, and S. Krinner, *Phys. Rev. Lett.* **134**, 120601 (2025).
- [22] L. Chen, S. P. Fors, Z. Yan, A. Ali, T. Abad, A. Osman, E. Moschandreou, B. Lienhard, S. Kosen, H.-X. Li *et al.*, [arXiv:2409.16748](https://arxiv.org/abs/2409.16748).
- [23] A. Kitaev, *Ann. Phys. (Amsterdam)* **303**, 2 (2003).
- [24] A. Eickbusch, M. McEwen, V. Sivak, A. Bourassa, J. Atalaya, J. Claes, D. Kafri, C. Gidney, C. W. Warren, J. Gross *et al.*, [arXiv:2412.14360](https://arxiv.org/abs/2412.14360).
- [25] D. Gao, D. Fan, C. Zha, J. Bei, G. Cai, J. Cai, S. Cao, F. Chen, J. Chen, K. Chen *et al.*, *Phys. Rev. Lett.* **134**, 090601 (2025).
- [26] T. Jiang, J. Cai, J. Huang, N. Zhou, Y. Zhang, J. Bei, G. Cai, S. Cao, F. Chen, J. Chen *et al.*, [arXiv:2505.01978](https://arxiv.org/abs/2505.01978).
- [27] See Supplemental Material at <http://link.aps.org/supplemental/10.1103/rqkg-dw31>, which includes Refs. [28–40], for detailed information and additional data.
- [28] F. Arute, K. Arya, R. Babbush, D. Bacon, J. C. Bardin, R. Barends, R. Biswas, S. Boixo, F. G. Brandao, D. A. Buell *et al.*, *Nature (London)* **574**, 505 (2019).
- [29] A. Bengtsson, A. Opremcak, M. Khezri, D. Sank, A. Bourassa, K. J. Satzinger, S. Hong, C. Erickson, B. J. Lester, K. C. Miao *et al.*, *Phys. Rev. Lett.* **132**, 100603 (2024).
- [30] J. Gambetta, A. Blais, D. I. Schuster, A. Wallraff, L. Frunzio, J. Majer, M. H. Devoret, S. M. Girvin, and R. J. Schoelkopf, *Phys. Rev. A* **74**, 042318 (2006).
- [31] 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, *Phys. Rev. Appl.* **10**, 034040 (2018).
- [32] C. A. Ryan, B. R. Johnson, J. M. Gambetta, J. M. Chow, M. P. da Silva, O. E. Dial, and T. A. Ohki, *Phys. Rev. A* **91**, 022118 (2015).
- [33] D. Sank, A. Opremcak, A. Bengtsson, M. Khezri, J. Chen, O. Naaman, and A. Korotkov, *Phys. Rev. Appl.* **23**, 024055 (2025).
- [34] C. C. Bultink, B. Tarasinski, N. Haandbæk, S. Poletto, N. Haider, D. J. Michalak, A. Bruno, and L. DiCarlo, *Appl. Phys. Lett.* **112**, 92601 (2018).
- [35] T. Gullion, D. B. Baker, and M. S. Conradi, *J. Magn. Reson.* (1969) **89**, 479 (1990).
- [36] J. Kelly, R. Barends, A. G. Fowler, A. Megrant, E. Jeffrey, T. C. White, D. Sank, J. Y. Mutus, B. Campbell, Y. Chen *et al.*, *Phys. Rev. A* **94**, 032321 (2016).
- [37] J. P. Bonilla Ataides, D. K. Tuckett, S. D. Bartlett, S. T. Flammia, and B. J. Brown, *Nat. Commun.* **12**, 2172 (2021).
- [38] O. Higgott, PyMatching: A Python package for quantum error-correction decoding with minimum-weight perfect matching, <https://github.com/oscarhiggott/PyMatching> (2025).
- [39] A. G. Fowler, [arXiv:1310.0863](https://arxiv.org/abs/1310.0863).
- [40] O. Higgott, BeliefMatching: An implementation of the belief-matching decoder (v0.1.0), <https://github.com/oscarhiggott/BeliefMatching> (2023).
- [41] F. Motzoi, J. M. Gambetta, P. Rebentrost, and F. K. Wilhelm, *Phys. Rev. Lett.* **103**, 110501 (2009).
- [42] S. Li, D. Fan, M. Gong, Y. Ye, X. Chen, Y. Wu, H. Guan, H. Deng, H. Rong, H.-L. Huang *et al.*, *Chin. Phys. Lett.* **39**, 030302 (2022).
- [43] V. Negîrneac, H. Ali, N. Muthusubramanian, F. Battistel, R. Sagastizabal, M. S. Moreira, J. F. Marques, W. J. Vlothuizen, M. Beekman, C. Zachariadis *et al.*, *Phys. Rev. Lett.* **126**, 220502 (2021).
- [44] D. Sank, Z. Chen, M. Khezri, J. Kelly, R. Barends, B. Campbell, Y. Chen, B. Chiaro, A. Dunsworth, A. Fowler *et al.*, *Phys. Rev. Lett.* **117**, 190503 (2016).
- [45] Y. Zhou, Z. Zhang, Z. Yin, S. Huai, X. Gu, X. Xu, J. Allcock, F. Liu, G. Xi, Q. Yu *et al.*, *Nat. Commun.* **12**, 5924 (2021).
- [46] F. Battistel, B. M. Varbanov, and B. M. Terhal, *PRX Quantum* **2**, 030314 (2021).
- [47] E. Magesan, J. M. Gambetta, B. R. Johnson, C. A. Ryan, J. M. Chow, S. T. Merkel, M. P. da Silva, G. A. Keefe, M. B. Rothwell, T. A. Ohki *et al.*, *Phys. Rev. Lett.* **109**, 080505 (2012).
- [48] O. Higgott and C. Gidney, *Quantum* **9**, 1600 (2025).
- [49] O. Higgott, T. C. Bohdanowicz, A. Kubica, S. T. Flammia, and E. T. Campbell, *Phys. Rev. X* **13**, 031007 (2023).
- [50] N. Lacroix, A. Bourassa, F. J. H. Heras, L. M. Zhang, J. Bausch, A. W. Senior, T. Edlich, N. Shutty, V. Sivak, A. Bengtsson *et al.*, *Nature (London)* **645**, 614 (2025).
- [51] S. Bravyi, A. W. Cross, J. M. Gambetta, D. Maslov, P. Rall, and T. J. Yoder, *Nature (London)* **627**, 778 (2024).