

Suppressing Spurious Transitions Using Spectrally Balanced Pulse

Ruixia Wang,^{1,*} Yaqing Feng,¹ Yujia Zhang,^{1,2,3} Jiayu Ding,¹ Boxi Li,^{4,5} Felix Motzoi,^{4,5} Yang Gao,^{1,2,3} Huikai Xu,¹ Zhen Yang,¹ Wuerkaixi Nuerbolati,¹ Haifeng Yu,¹ Weijie Sun,^{1,†} and Fei Yan^{1,‡}

¹*Beijing Key Laboratory of Fault-Tolerant Quantum Computing,
Beijing Academy of Quantum Information Sciences, Beijing 100193, China*

²*Beijing National Laboratory for Condensed Matter Physics and
Institute of Physics, Chinese Academy of Sciences, Beijing 100190, China*

³*University of Chinese Academy of Sciences, Beijing 101408, China*

⁴*Forschungszentrum Jülich, Institute of Quantum Control (PGI-8), D-52425 Jülich, Germany*

⁵*Institute for Theoretical Physics, University of Cologne, D-50937 Cologne, Germany*

 (Received 14 February 2025; revised 24 April 2025; accepted 23 September 2025; published 14 October 2025)

Achieving precise control over quantum systems presents a significant challenge, especially in many-body setups, where residual couplings and unintended transitions undermine the accuracy of quantum operations. In superconducting qubits, parasitic interactions—both between distant qubits and with spurious two-level systems—can severely limit the performance of quantum gates. In this Letter, we introduce a pulse-shaping technique that uses spectrally balanced microwave pulses to suppress undesired transitions. Experimental results demonstrate an order-of-magnitude reduction in spurious excitations between weakly detuned qubits as well as a substantial decrease in single-qubit gate errors caused by a strongly coupled two-level defect over a broad frequency range. Our method provides a simple yet powerful solution for mitigating adverse effects from parasitic couplings, enhancing quantum gate fidelity. The pulse-shaping technique can be readily adapted to various physical systems.

DOI: [10.1103/h4xf-vq2l](https://doi.org/10.1103/h4xf-vq2l)

Introduction—Quantum information processing technologies have made significant progress, achieving coherent control over systems with about 100 qubits [1–6]. In superconducting quantum processors, qubit-qubit interactions are typically mediated by engineered coupling elements, such as coplanar capacitors. However, due to the long-range nature of electromagnetic interactions and design constraints, parasitic or residual couplings can exist between qubits that are intended to be uncoupled. These parasitic interactions, known as quantum crosstalk [7], pose a challenge to the execution of independent operations, as they can interfere with transitions close in frequency. This phenomenon degrades the performance of quantum operations, limiting the scalability of quantum computing systems.

Compounding this issue is the presence of uncontrolled microscopic degrees of freedom, often referred to as two-level defects or two-level systems (TLSs), which are ubiquitous in many quantum platforms [8]. When frequencies of TLSs come close to those of the qubits, they can significantly interfere with qubit operations. Excitations into a long-lived TLS can be particularly harmful as they can accumulate over time and deteriorate subsequent operations, causing correlated or non-Markovian errors

[9]. Without reset capability, this may be a potential threat to quantum error correction [10]. Moving the qubit frequency away from these TLSs can reduce their impact. However, this approach is not directly available for fixed-frequency qubits. Even for tunable qubits, the available frequency options are often severely constrained by the presence of multiple TLSs, complicating the calibration of large-scale processors [11].

Various techniques have been developed to address these challenges [12–15]. The derivative removal by adiabatic gate (DRAG) method [16–19], in particular, was a pulse-shaping technique designed to suppress leakage transitions to higher-energy states during single-qubit operations. By adding the derivative of the original pulse envelope to the quadrature component, DRAG removes unwanted diabatic leakage transitions, which, in the weak drive regime, show up as a spectral hole at a frequency determined by the prefactor of the derivative. The versatility of DRAG has been investigated in a variety of scenarios, including three-level lambda systems [20], cross-resonance gates [21–23], and frequency-crowded multilevel systems [24]. However, traditional DRAG corrections face challenges when dealing with weakly detuned transitions, as they struggle to effectively remove spectral components close to the target transition [22].

In this Letter, we propose and experimentally demonstrate a robust approach to mitigating undesired transitions caused by quantum crosstalk during single-qubit

*Contact author: wangrx@baqis.ac.cn

†Contact author: sunwj@baqis.ac.cn

‡Contact author: yanfei@baqis.ac.cn

operations. Unlike conventional DRAG techniques, our method employs the dual-DRAG protocol, which creates symmetric spectral holes around the target transition frequency. This approach significantly reduces off-resonance effects during pulses and enables the pulse calibration to suppress unwanted transitions that are only slightly detuned. Residual off-resonance effects are further corrected using compensatory virtual-Z (VZ) gates. Experimental results validate the technique, demonstrating an order-of-magnitude suppression of crosstalk-induced excitations between two superconducting qubits detuned by approximately 40 MHz with a 25-ns pulse across various coupling strengths. Furthermore, we show its effectiveness in reducing excitations from a strongly coupled TLS and enhancing single-qubit gate fidelity across a broad frequency range, extending down to 20 MHz.

Quantum crosstalk—Parasitic couplings are a common challenge in practical quantum hardware, often arising between qubits that are not intentionally connected. For example, in solid-state devices like superconducting quantum processors, qubits are typically designed for nearest-neighbor connections. However, the long-range nature of electromagnetic interactions can lead to unintended couplings between physically distant qubits, as shown in Fig. 1(a). These interactions can also occur through shared

modes, such as chip or box modes. Additionally, unwanted couplings can exist between qubits and spurious quantum systems in the environment, such as TLSs. These uncontrollable systems present a major obstacle in the development and calibration of state-of-the-art quantum hardware.

Consider a parasitic exchange-type coupling g between Q_0 and Q_1 (or a TLS). Because of the hybridization of the $|ge\rangle$ and $|eg\rangle$ states, driving one qubit inevitably induces a cross-driving effect on the other qubit or TLS, but with a reduced amplitude of $g\Omega/\Delta_0$, where Ω is the original drive amplitude applied to Q_0 and $\Delta_0 = \omega_{ge} - \omega_{eg}$ is the detuning between the unwanted and target transitions, as illustrated in Fig. 1(b). Natural spectral broadening caused by finite pulse widths can lead to unwanted transitions, such as $|gg\rangle \leftrightarrow |ge\rangle$, which degrade the performance of quantum operations. This crosstalk phenomenon is inherently quantum mechanical; the crosstalk Hamiltonian is of the ZX type, exhibiting opposite signs depending on the state of Q_0 . This contrasts with classical signal crosstalk, where the Hamiltonian is of the IX type. As a result, quantum crosstalk cannot be simply corrected by applying a cancellation signal [25]. Interestingly, this crosstalk effect can be leveraged to construct a type of two-qubit entangling gate, the cross-resonance gate [26–28].

Dual-DRAG—A straightforward solution to suppress an unwanted transition is to remove the spectral components at that transition frequency. A famous example is the DRAG method, which adds a first-order derivative of the pulse envelope as a quadrature component, with an amplitude factor $1/\Delta$, where Δ is the DRAG parameter in angular frequency units. An example is illustrated in Fig. 1(c). This creates a spectral hole at frequency Δ relative to the drive frequency, as depicted in Fig. 1(d), effectively reducing spectral components near the qubit’s $|e\rangle$ - $|f\rangle$ transition—typically, 200–300 MHz below the $|g\rangle$ - $|e\rangle$ transition for transmon qubits—and suppressing leakage transitions.

Avoiding weakly detuned transitions during single-qubit operations, which are only a few tens of megahertz away, has proven to be challenging [22]. For instance, as shown in Figs. 1(c) and 1(d), a 25-ns pulse with a single DRAG correction of $\Delta/2\pi = 40$ MHz shifts the drive frequency by approximately 30 MHz. When the DRAG-induced off-resonance effect is this large, the perturbative treatment in the DRAG theory breaks down, often causing the conventional mitigation approach which introduces a constant drive frequency detuning η [18] to fail.

To address this issue, we adopt the recursive DRAG scheme, where DRAG corrections are applied sequentially according to

$$\Omega^{(n)} = \Omega^{(n-1)} - i \frac{\dot{\Omega}^{(n-1)}}{\Delta_n}, \quad (1)$$

using a set of DRAG parameters $\{\Delta_1, \Delta_2, \dots, \Delta_n\}$ [23,29]. Here, $\Omega^{(0)}$ represents the original pulse shape without DRAG. The final pulse $\Omega^{(n)}$ generates multiple spectral

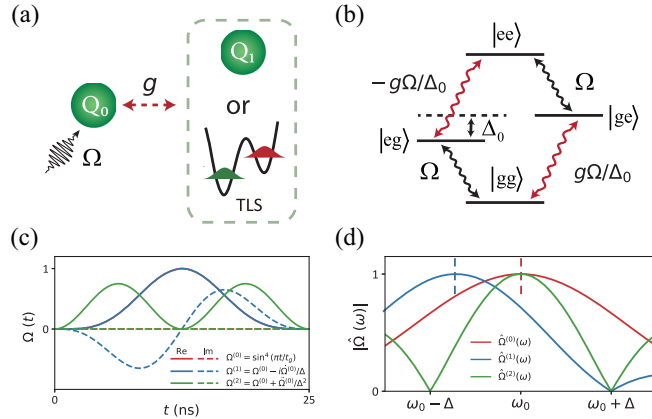


FIG. 1. (a) Schematic diagram illustrating the unintended coupling (g) that can occur between two qubits or between a qubit and a spurious TLS. (b) Energy-level diagram for the combined system of Q_0 and Q_1 (or TLS). Because of state dressing, driving Q_0 with an amplitude Ω also cross-drives Q_1 with a reduced amplitude of $\pm g\Omega/\Delta_0$ in the weak coupling limit ($g \ll |\Delta_0|$). Cross-driving can induce unwanted transitions in Q_1 when using a pulse with finite width. (c) Time-domain pulse profiles and (d) their normalized Fourier spectra for a 25-ns pulse, showing its original sine-to-the-fourth-power form ($\Omega^{(0)}$), the first-order DRAG-corrected pulse ($\Omega^{(1)}$), and the spectrally balanced pulse with dual-DRAG corrections ($\Omega^{(2)}$). Note that, in the time domain, the real part of $\Omega^{(0)}$ overlaps with that of $\Omega^{(1)}$, while the imaginary part of $\Omega^{(0)}$ overlaps with that of $\Omega^{(2)}$. The vertical dashed lines in (d) indicate the corresponding peak positions. In this case, the frequency to block is assumed to be $\Delta/2\pi = 40$ MHz.

holes at the specified frequencies. By two successive DRAG applications at $\pm\Delta$, we obtain a pulse envelope that involves a second derivative:

$$\Omega^{(2)} = \Omega^{(0)} + \frac{\ddot{\Omega}^{(0)}}{\Delta^2}. \quad (2)$$

As shown in Fig. 1(d), the spectrum $\hat{\Omega}^{(2)}(\omega)$ remains centered at ω_0 , while spectral components at $\omega_0 \pm \Delta$ are eliminated. This dual-DRAG application at mirrored frequencies suppresses weakly detuned transitions and significantly reduces drive-induced off-resonance [30]. In the example presented, we choose the original pulse shape to be a sine raised to the fourth power. This ensures that the first three derivatives vanish at the pulse's beginning and end, thereby avoiding sharp edges or the need for truncation.

Notably, Ref. [19] reports remarkable leakage suppression through their implementation of multiple derivative techniques. Their method, while not explicitly discussed, effectively utilizes an in-phase component analogous to Eq. (2), applying successive DRAG corrections at $\pm\alpha$ followed by an additional correction at α . Crucially, this arrangement naturally improves spectral symmetry—a fundamental property that, as we have identified, plays a pivotal role in mitigating DRAG-induced off-resonance effects and enabling their successful high-order derivative suppression.

Correcting off-resonance error using virtual-Z compensation—Although the dual-DRAG method substantially reduces the deviation of spectral weighting, system nonidealities such as the presence of higher energy levels and pulse distortions can lead to residual off-resonance effects. To address this, we adopt a postcorrection approach that compensates for the overall effect of an off-resonant pulse by appending a VZ gate. We validate this protocol on a device similar to that described in Ref. [31] (see Supplemental Material [30] for more information). For all experiments in this Letter unless specified otherwise, we constantly apply one DRAG correction at the qubit anharmonicity $\alpha/2\pi \approx -190$ MHz to prohibit leakage transition to the $|f\rangle$ state.

Throughout our Letter, we employ the U3 decomposition rule to synthesize arbitrary single-qubit operations using two $\sqrt{X}$ gates interspersed with VZ gates [32]. Consequently, it suffices to calibrate only the $\sqrt{X}$ gate. It can be shown that any rotation R , which maps the qubit from its ground state to the equator (XY plane) of the Bloch sphere, plus a virtual Z or phase rotation, can be equivalently used as a $\sqrt{X}$ gate (see Supplemental Material [30] and Ref. [32]). Thus, calibrating such an R operation suffices to synthesize any single-qubit gate with the aid of VZ gates.

To calibrate the R gate for a given pulse shape, duration, carrier frequency, and DRAG parameters, we scan the drive

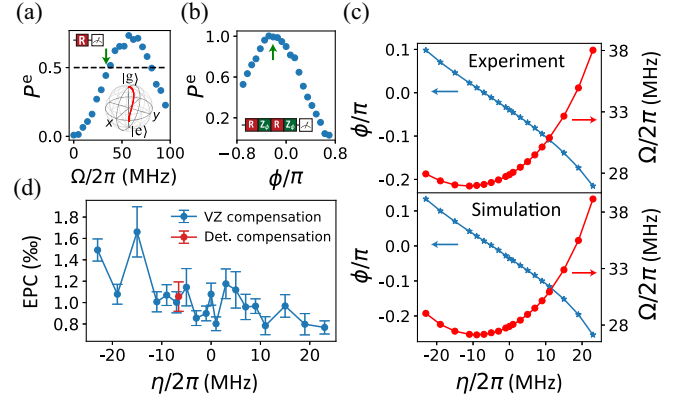


FIG. 2. (a) Calibrating the pulse amplitude of the R gate ($t_g = 25$ ns) which drives the qubit from its ground state to an equator state, as indicated by the arrow. The drive detuning is set at $\eta/2\pi = 23$ MHz. The inset plots are the circuit and a sketch of the state evolution on the Bloch sphere. (b) Calibrating the compensating virtual- Z phase after the R gate, which yields an $\sqrt{X}$ gate. By applying the combination twice, the qubit is rotated from the ground state to the excited state, as indicated by the arrow. The inset is the circuit. (c) Experimental results (top panel) showing the calibrated drive amplitude and virtual- Z phase for different drive detunings, compared with numerical simulations (bottom panel). (d) Measured error per Clifford (EPC) from randomized benchmarking using pulses with virtual- Z compensation (blue) for different drive detunings. Results using detuning compensation (without the appended virtual- Z phase, shown in red) are also included for comparison.

amplitude until the excited-state population of the qubit (initialized in its ground state) reaches 0.5 [Fig. 2(a)]. We then concatenate two R gates with VZ gates to determine the necessary VZ phases for maximal excitation [Fig. 2(b)]. After fine-tuning these parameters using standard pulse techniques [30], we validate the calibrated parameters Ω (drive amplitude) and ϕ (VZ phase) under intentional detunings, with results closely matching numerical simulations [Fig. 2(c)]. Using these VZ-compensated $\sqrt{X}$ gates, we perform randomized benchmarking (RB) [33]. The errors per Clifford (EPC), shown in Fig. 2(d), demonstrate consistent performance over a detuning range of ± 20 MHz, showing that the VZ compensation method is suited for more general scenarios.

Suppressing spurious transitions between qubits—Next, we validate the dual-DRAG method's effectiveness in suppressing unwanted transitions due to quantum crosstalk. We use two frequency-tunable transmon qubits, with a tunable coupler (another transmon) adjusting the effective qubit-qubit coupling strength g via the coupler frequency ω_c [Fig. 3(a)]. At first, we set the qubit-qubit detuning $\Delta_0/2\pi = 45$ MHz and the coupling $g/2\pi = 1$ MHz.

For efficient detection of weak transitions ($g\Omega/\Delta_0 < \Delta_0$), we employ a protocol similar to the interference error filter [34,35] [Fig. 3(b)]. An even number ($2N$) of X gates are applied to the control qubit Q_0 , keeping it in the ground

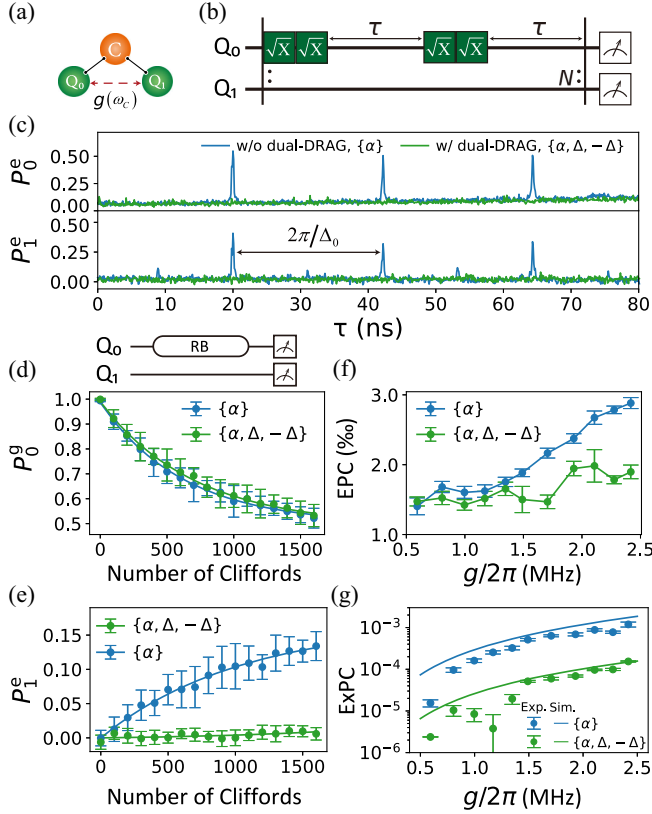


FIG. 3. (a) Experimental setup consisting of two transmon qubits and a tunable coupler, which is also a transmon qubit. The coupler is used to modulate the effective coupling strength between the qubits by adjusting its frequency. (b) Pulse sequence designed for detecting spurious transitions in a nearby qubit during single-qubit gate operations. Repeated π gates, each consisting of two $\pi/2$ pulses and spaced by a waiting time τ , are applied to Q_0 . (c) Measured excited-state populations of Q_0 (top panel) and Q_1 (bottom panel) using the detection sequence, as a function of the waiting time τ . Results are shown with and without dual-DRAG. Both cases include DRAG correction at the leakage transition to the second excited state $|2\rangle$. The number of π gate pairs is $N = 50$. Here, $g/2\pi = 1$ MHz, $\Delta_0/2\pi = 45$ MHz, and $t_g = 25$ ns. (d) Randomized benchmarking results for Q_0 , comparing cases with and without dual-DRAG. (e) Simultaneously monitored excited-state populations of Q_1 . The excitation rates per Clifford gate (ExPC) are $(0.8 \pm 0.3) \times 10^{-5}$ with dual-DRAG and $(1.6 \pm 0.2) \times 10^{-4}$ without dual-DRAG. (f) Error per Clifford for different coupling strengths, with a fixed detuning of $\Delta_0/2\pi = 45$ MHz. (g) Excitation per Clifford for varying coupling strengths. The solid lines are the simulated ExPC averaged over 24 single-qubit Cliffords implemented with the U3 decomposition. The excitation rates are proportional to g^2 (see Supplemental Material [30] for details).

state. Because of nonzero coupling, the spectator qubit Q_1 is coherently excited by these pulses. A uniform waiting time τ is added after each π gate, during which the wave function amplitude of Q_1 's excited state (e.g., $|ge\rangle$) accumulates a relative phase $e^{-i\Delta_0\tau}$. At specific τ , the transition amplitudes interfere constructively, amplifying Q_1 's excitation.

Figure 3(c) shows the measured total populations of all excited states (P_i^e , $i = 0, 1$), including $|e\rangle$ and $|f\rangle$, for both qubits as a function of τ . Without dual-DRAG correction, strong periodic peaks appear on both qubits at intervals of ~ 22.2 ns, matching $2\pi/\Delta_0$ and confirming our expectation. Since X gates are applied in pairs, transitions to the excited state of Q_1 —whether through $|gg\rangle \rightarrow |ge\rangle$ or $|eg\rangle \rightarrow |ee\rangle$ —result in the $|ee\rangle$ state in the end. Additionally, smaller secondary peaks on Q_1 , halfway between primary peaks and sharing the same periodicity, are attributed to residual IX interaction arising from both quantum and classical crosstalk [30].

We repeat measurements with dual-DRAG correction at $\pm\Delta$. The parameter Δ is initially set to Δ_0 and fine-tuned by minimizing excitation peaks. Using three DRAG parameters $\{\alpha, \Delta, -\Delta\}$, all prominent excitation peaks vanish [Fig. 3(c)], demonstrating effective suppression. Randomized benchmarking on Q_0 [Figs. 3(d) and 3(e)] shows a slight error reduction (error per Clifford: $1.60 \times 10^{-3} \rightarrow 1.42 \times 10^{-3}$). Simultaneously, Q_1 's excitation rate per Clifford (ExPC) drops by an order of magnitude from $(1.6 \pm 0.2) \times 10^{-4}$ to $(0.8 \pm 0.3) \times 10^{-5}$. Although the effect in EPC does not seem to be significant, the induced correlated excitations, in particular between distant qubits, complicate the error model and may severely affect error correction codes.

We further test the method across coupling strengths using the tunable coupler. The improvement in gate error rates using dual-DRAG becomes more pronounced at larger g [Fig. 3(f)]. This trend aligns with the measured spurious excitation rates, which scale quadratically with g as expected from the cross-driving amplitude that scales with g/Δ_0 . At large g , these spurious excitations dominate the gate errors, yet dual-DRAG consistently suppresses them by an order of magnitude across the whole range [Fig. 3(g)]. These results demonstrate the effectiveness of spectrally balanced DRAG in suppressing weakly detuned transitions and enhancing gate fidelities.

Applications to the qubit-TLS system—We extend our method to suppress spurious transitions in systems where Q_0 is coupled to a parasitic TLS. Such TLSs with long relaxation times are particularly problematic, as their unintended excitations can accumulate over time and persistently degrade subsequent gate operations. The specific TLS studied here exhibits a strong coupling ($g/2\pi = 5$ MHz) to a transmon qubit, as evidenced by the avoided crossing in the measured qubit spectrum shown in Fig. 4(a).

To characterize the dependence on detuning of our method, we calibrate the $\sqrt{X}$ gate at varying qubit-TLS detunings Δ_0 using the same protocol developed for coupled qubits. Figure 4(b) shows RB results comparing dual-DRAG corrections (applied at $\pm\Delta$) to uncorrected pulses across $20 \text{ MHz} \leq \Delta_0/2\pi \leq 85 \text{ MHz}$. For $\Delta_0/2\pi < 20 \text{ MHz}$, spectral overlap between the drive frequency and

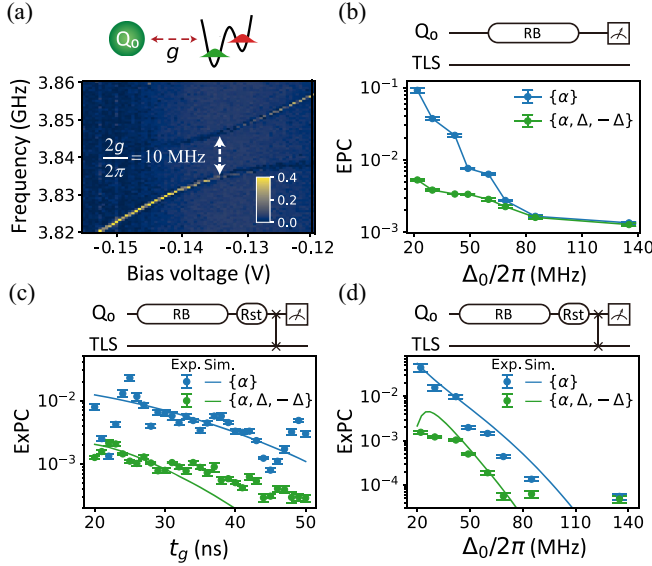


FIG. 4. (a) Spectroscopy of a tunable transmon qubit coupled to a spurious TLS, showing an avoided crossing of approximately 10 MHz. The horizontal axis represents the voltage of the flux bias pulse applied to the coupler qubit. (b) Error per Clifford of the qubit as a function of detuning between the qubit and the TLS. Here, $t_g = 25$ ns. (c) Excitation per Clifford of the TLS for different gate times at $\Delta_0/2\pi = 42$ MHz. The solid lines are twice the simulated $\sqrt{X}$ gate errors. (d) Excitation per Clifford of the TLS for varying detunings. The solid lines are twice the simulated $\sqrt{X}$ gate errors.

the TLS transition hinders the convergence of parameters during calibration, and the dual-DRAG correction overcuts the spectral components around the target frequency, making the drive ineffective. At $\Delta_0/2\pi > 85$ MHz, TLS-induced gate errors diminish to a negligible level.

We further quantify TLS excitation rates across gate time t_g and detuning Δ_0 [Figs. 4(c) and 4(d)]. To measure TLS excitations, we (i) reset the qubit to its ground state post-RB, (ii) apply an iSWAP gate to transfer TLS excitations to the qubit, and (iii) perform qubit readout. In $20 \text{ ns} \leq t_g \leq 50 \text{ ns}$, dual-DRAG reduces TLS excitation rates by an order of magnitude compared to uncorrected pulses. Numerical simulations of single-pulse dynamics reproduce the observed decreasing trend, though deviations may arise from interference between pulses and spectral smoothing due to the presence of dephasing noise.

The measured excitation rates for $20 \text{ MHz} \leq \Delta_0/2\pi \leq 85 \text{ MHz}$ show a comparable improvement from the use of dual-DRAG. The rapidly enhanced suppression with increasing Δ_0 results from the diminishing spectral weight at larger Δ_0 and the reduced cross-drive amplitude that scales with g/Δ_0 .

Discussion—We have demonstrated a spectrally balanced pulse-shaping technique that suppresses spurious transitions in superconducting qubits. Our method reduces crosstalk errors by an order of magnitude and mitigates

single-qubit gate errors from strongly coupled TLS. By engineering symmetric spectral holes, it effectively suppresses weakly detuned transitions by mitigating the off-resonance effect during pulses. The technique integrates seamlessly with virtual-Z gates, further improving gate fidelity.

This approach directly addresses key challenges in a frequency-crowded quantum processor: parasitic qubit couplings and TLS-induced errors, both of which hinder frequency allocation. By mitigating these effects, our Letter can enhance frequency planning flexibility and provide a framework for modeling microwave gate errors in complex environments.

We note that, while dual-DRAG pulses demonstrate strong performance for small-scale systems (as experimentally validated), their large-scale efficacy requires careful frequency arrangement to avoid extreme conditions with multiple near-resonant TLS (see Supplemental Material [30]). Future work should systematically investigate these scaling properties. However, this scheme can be effectively combined with frequency-tuning approaches that isolate TLS from qubits [36–38] to further enhance gate performance, providing a comprehensive solution to the TLS problem.

The hardware-agnostic nature of the dual-DRAG protocol makes it broadly applicable to other quantum architectures, including superconducting fluxonium qubits, semiconductor quantum dots, NV centers, and trapped ions, positioning it as a versatile tool for high-fidelity quantum control [20,39–42].

Note added—Recently, we became aware of a concurrent work [43] that independently develops a related pulse-shaping technique using symmetrically filtered spectra to address classical crosstalk.

Acknowledgments—The authors thank Yu He, Yasunobu Nakamura, and Peng Zhao for fruitful discussions. This work was supported by Beijing Natural Science Foundation (Grant No. JQ25014), National Natural Science Foundation of China (Grants No. 12404558, No. 12322413, No. 92365206, and No. 92476206), and Innovation Program for Quantum Science and Technology (Grant No. 2021ZD0301802).

Data availability—The data that support the findings of this article are available from the authors upon reasonable request.

- [1] Y. Kim, A. Eddins, S. Anand, K. X. Wei, E. Van Den Berg, S. Rosenblatt, H. Nayfeh, Y. Wu, M. Zaletel, K. Temme, and A. Kandala, Evidence for the utility of quantum computing before fault tolerance, *Nature (London)* **618**, 500 (2023).
- [2] R. Acharya *et al.*, Quantum error correction below the surface code threshold, *Nature (London)* **638**, 920 (2024).

- [3] S. Cao *et al.*, Generation of genuine entanglement up to 51 superconducting qubits, *Nature (London)* **619**, 738 (2023).
- [4] S. Xu *et al.*, Digital simulation of projective non-Abelian anyons with 68 superconducting qubits, *Chin. Phys. Lett.* **40**, 060301 (2023).
- [5] M. Iqbal, N. Tantivasadakarn, R. Verresen, S. L. Campbell, J. M. Dreiling, C. Figgatt, J. P. Gaebler, J. Johansen, M. Mills, S. A. Moses, J. M. Pino, A. Ransford, M. Rowe, P. Siegfried, R. P. Stutz, M. Foss-Feig, A. Vishwanath, and H. Dreyer, Non-Abelian topological order and anyons on a trapped-ion processor, *Nature (London)* **626**, 505 (2024).
- [6] D. Bluvstein, S. J. Evered, A. A. Geim, S. H. Li, H. Zhou, T. Manovitz, S. Ebadi, M. Cain, M. Kalinowski, D. Hangleiter, J. P. Bonilla Ataides, N. Maskara, I. Cong, X. Gao, P. Sales Rodriguez, T. Karolyshyn, G. Semeghini, M. J. Gullans, M. Greiner, V. Vuletić, and M. D. Lukin, Logical quantum processor based on reconfigurable atom arrays, *Nature (London)* **626**, 58 (2024).
- [7] M. Sarovar, T. Proctor, K. Rudinger, K. Young, E. Nielsen, and R. Blume-Kohout, Detecting crosstalk errors in quantum information processors, *Quantum* **4**, 321 (2020).
- [8] C. Müller, J. H. Cole, and J. Lisenfeld, Towards understanding two-level-systems in amorphous solids: Insights from quantum circuits, *Rep. Prog. Phys.* **82**, 124501 (2019).
- [9] M. Y. Niu, V. Smelyanskiy, P. Klimov, S. Boixo, R. Barends, J. Kelly, Y. Chen, K. Arya, B. Burkett, D. Bacon *et al.*, Learning non-Markovian quantum noise from moiré-enhanced swap spectroscopy with deep evolutionary algorithm, [arXiv:1912.04368](https://arxiv.org/abs/1912.04368).
- [10] J. Ghosh, A. G. Fowler, J. M. Martinis, and M. R. Geller, Understanding the effects of leakage in superconducting quantum-error-detection circuits, *Phys. Rev. A* **88**, 062329 (2013).
- [11] P. V. Klimov *et al.*, Optimizing quantum gates towards the scale of logical qubits, *Nat. Commun.* **15**, 2442 (2024).
- [12] 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).
- [13] Y. Ding, P. Gokhale, S. F. Lin, R. Rines, T. Propson, and F. T. Chong, Systematic crosstalk mitigation for superconducting qubits via frequency-aware compilation, in *Proceedings of the 2020 53rd Annual IEEE/ACM International Symposium on Microarchitecture (MICRO)* (IEEE, New York, 2020), pp. 201–214.
- [14] P. Murali, D. C. McKay, M. Martonosi, and A. Javadi-Abhari, Software mitigation of crosstalk on noisy intermediate-scale quantum computers, in *Proceedings of the Twenty-Fifth International Conference on Architectural Support for Programming Languages and Operating Systems* (2020), pp. 1001–1016.
- [15] P. Zhao, Mitigation of quantum crosstalk in cross-resonance-based qubit architectures, *Phys. Rev. Appl.* **20**, 054033 (2023).
- [16] F. Motzoi, J. M. Gambetta, P. Rebentrost, and F. K. Wilhelm, Simple pulses for elimination of leakage in weakly nonlinear qubits, *Phys. Rev. Lett.* **103**, 110501 (2009).
- [17] J. M. Gambetta, F. Motzoi, S. Merkel, and F. K. Wilhelm, Analytic control methods for high-fidelity unitary operations in a weakly nonlinear oscillator, *Phys. Rev. A* **83**, 012308 (2011).
- [18] Z. Chen *et al.*, Measuring and suppressing quantum state leakage in a superconducting qubit, *Phys. Rev. Lett.* **116**, 020501 (2016).
- [19] E. Hyppä, A. Vepsäläinen, M. Papič, C. F. Chan, S. Inel, A. Landra, W. Liu, J. Luus, F. Marxer, C. Ockeloen-Korppi, S. Orbell, B. Tarasinski, and J. Heinsoo, Reducing leakage of single-qubit gates for superconducting quantum processors using analytical control pulse envelopes, *PRX Quantum* **5**, 030353 (2024).
- [20] A. Vezvaei, E. Takou, P. Hilaire, M. F. Doty, and S. E. Economou, Avoiding leakage and errors caused by unwanted transitions in lambda systems, *PRX Quantum* **4**, 030312 (2023).
- [21] M. Malekakhlagh and E. Magesan, Mitigating off-resonant error in the cross-resonance gate, *Phys. Rev. A* **105**, 012602 (2022).
- [22] K. X. Wei, E. Pritchett, D. M. Zajac, D. C. McKay, and S. Merkel, Characterizing non-Markovian off-resonant errors in quantum gates, *Phys. Rev. Appl.* **21**, 024018 (2024).
- [23] B. Li, T. Calarco, and F. Motzoi, Experimental error suppression in cross-resonance gates via multi-derivative pulse shaping, *npj Quantum Inf.* **10**, 66 (2024).
- [24] L. Theis, F. Motzoi, and F. Wilhelm, Simultaneous gates in frequency-crowded multilevel systems using fast, robust, analytic control shapes, *Phys. Rev. A* **93**, 012324 (2016).
- [25] W. Nuerbolati, Z. Han, J. Chu, Y. Zhou, X. Tan, Y. Yu, S. Liu, and F. Yan, Canceling microwave crosstalk with fixed-frequency qubits, *Appl. Phys. Lett.* **120** (2022).
- [26] C. Rigetti and M. Devoret, Fully microwave-tunable universal gates in superconducting qubits with linear couplings and fixed transition frequencies, *Phys. Rev. B* **81**, 134507 (2010).
- [27] P. De Groot, J. Lisenfeld, R. Schouten, S. Ashhab, A. Lupaşcu, C. Harmans, and J. Mooij, Selective darkening of degenerate transitions demonstrated with two superconducting quantum bits, *Nat. Phys.* **6**, 763 (2010).
- [28] V. Tripathi, M. Khezri, and A. N. Korotkov, Operation and intrinsic error budget of a two-qubit cross-resonance gate, *Phys. Rev. A* **100**, 012301 (2019).
- [29] F. Motzoi and F. K. Wilhelm, Improving frequency selection of driven pulses using derivative-based transition suppression, *Phys. Rev. A* **88**, 062318 (2013).
- [30] See Supplemental Material at <http://link.aps.org/supplemental/10.1103/h4xf-vq2l> for the supporting information of the experiments and simulations.
- [31] Z. Chen, W. Liu, Y. Ma, W. Sun, R. Wang, H. Wang, H. Xu, G. Xue, H. Yan, Z. Yang, J. Ding, Y. Gao, F. Li, Y. Zhang, Z. Zhang, Y. Jin, H. Yu, J. Chen, and F. Yan, Efficient implementation of arbitrary two-qubit gates via unified control, *Nat. Phys.* **21**, 1489 (2025).
- [32] 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).
- [33] E. Magesan, J. M. Gambetta, and J. Emerson, Scalable and robust randomized benchmarking of quantum processes, *Phys. Rev. Lett.* **106**, 180504 (2011).
- [34] E. Lucero, M. Hofheinz, M. Ansmann, R. C. Bialczak, N. Katz, M. Neeley, A. D. O'Connell, H. Wang, A. N. Cleland,

- and J. M. Martinis, High-fidelity gates in a single Josephson qubit, *Phys. Rev. Lett.* **100**, 247001 (2008).
- [35] R. Acharya *et al.*, Suppressing quantum errors by scaling a surface code logical qubit, *Nature (London)* **614**, 676 (2023).
- [36] A. Dane, K. Balakrishnan, B. Wacaser, L.-W. Hung, H. J. Mamin, D. Rugar, R. M. Shelby, C. Murray, K. Rodbell, and J. Sleight, Performance stabilization of high-coherence superconducting qubits, [arXiv:2503.12514](#).
- [37] L. Chen, K.-H. Lee, C.-H. Liu, B. Marinelli, R. K. Naik, Z. Kang, N. Goss, H. Kim, D. I. Santiago, and I. Siddiqi, Scalable and site-specific frequency tuning of two-level system defects in superconducting qubit arrays, [arXiv:2503.04702](#).
- [38] Z. T. Wang, P. Zhao, Z. H. Yang, Y. Tian, H. F. Yu, and S. P. Zhao, Escaping detrimental interactions with microwave-dressed transmon qubits, *Chin. Phys. Lett.* **40**, 070304 (2023).
- [39] I. Heinz and G. Burkard, Crosstalk analysis for simultaneously driven two-qubit gates in spin qubit arrays, *Phys. Rev. B* **105**, 085414 (2022).
- [40] C. Fang, Y. Wang, S. Huang, K. R. Brown, and J. Kim, Crosstalk suppression in individually addressed two-qubit gates in a trapped-ion quantum computer, *Phys. Rev. Lett.* **129**, 240504 (2022).
- [41] L. S. Theis, F. Motzoi, F. K. Wilhelm, and M. Saffman, High fidelity rydberg-blockade entangling gate using shaped, analytic pulses, *Phys. Rev. A* **94**, 032306 (2016).
- [42] K. Takeda, J. Yoneda, T. Otsuka, T. Nakajima, M. R. Delbecq, G. Allison, Y. Hoshi, N. Usami, K. M. Itoh, S. Oda, T. Kadera, and S. Tarucha, Optimized electrical control of a Si/SiGe spin qubit in the presence of an induced frequency shift, *npj Quantum Inf.* **4**, 54 (2018).
- [43] R. Matsuda, R. Ohira, T. Sumida, H. Shiomi, A. Machino, S. Morisaka, K. Koike, T. Miyoshi, Y. Kurimoto, Y. Sugita, Y. Ito, Y. Suzuki, P. A. Spring, S. Wang, S. Tamate, Y. Tabuchi, Y. Nakamura, K. Ogawa, and M. Negoro, Selective excitation of superconducting qubits with a shared control line through pulse shaping, [arXiv:2501.10710](#).