
Supplementary information

Probing entanglement in a 2D hard-core Bose–Hubbard lattice

In the format provided by the
authors and unedited

Supplementary material for “Probing entanglement in a 2D hard-core Bose-Hubbard lattice”

Amir H. Karamlou,^{1,2,*} Ilan T. Rosen,¹ Sarah E. Muschinske,^{1,2} Cora N. Barrett,^{1,3}
Agustin Di Paolo,¹ Leon Ding,^{1,4} Patrick M. Harrington,¹ Max Hays,¹ Rabindra Das,⁵
David K. Kim,⁵ Bethany M. Niedzielski,⁵ Meghan Schuldt,⁵ Kyle Serniak,^{1,5} Mollie E. Schwartz,⁵
Jonilyn L. Yoder,⁵ Simon Gustavsson,¹ Yariv Yanay,⁶ Jeffrey A. Grover,¹ and William D. Oliver^{1,2,4,5,†}

¹*Research Laboratory of Electronics, Massachusetts Institute of Technology, Cambridge, MA 02139, USA*

²*Department of Electrical Engineering and Computer Science,*

Massachusetts Institute of Technology, Cambridge, MA 02139, USA

³*Department of Physics, Wellesley College, Wellesley, MA 02481, USA*

⁴*Department of Physics, Massachusetts Institute of Technology, Cambridge, MA 02139, USA*

⁵*MIT Lincoln Laboratory, Lexington, MA 02421, USA*

⁶*Laboratory for Physical Sciences, College Park, MD 20740, USA*

(Dated: February 27, 2024)

CONTENTS

S1. Experimental setup	S2
S2. Sample	S4
S3. Control and readout calibration	S5
A. Flux crosstalk calibration	S5
B. Time-domain flux pulse shaping and transient calibration	S8
C. Pulse alignment	S9
D. Readout optimization	S10
E. State preparation and measurement (SPAM) error mitigation	S12
F. Idling frequency layout	S12
S4. Hamiltonian Characterization	S13
A. Particle exchange strengths	S13
B. Drive coupling amplitude	S13
C. Drive coupling phase	S14
S5. Hard-core Bose-Hubbard model energy spectrum	S16
S6. Coherent-like states across the spectrum	S17
S7. Impact of the drive coupling phase on experiments	S18
S8. Measurement of global entanglement	S18
S9. Tomography subsystems for studying the entanglement behavior	S19
S10. Extracting s_V and s_A	S22
S11. Extracting the entanglement scaling behavior in larger lattices	S23
S12. Schmidt coefficients scaling	S24
S13. Estimating the impact of decoherence on measured system entropy	S25

* karamlou@mit.edu

† william.oliver@mit.edu

S14. Time dynamics simulations	S27
S15. Coherent-like states in 1D	S28
S16. Extended data	S30
References	S33

S1. EXPERIMENTAL SETUP

We house our superconducting qubit array within a BlueFors XLD-600 dilution refrigerator. In our setup, the mixing chamber (MXC) base temperature reaches as low as 14mK, however, we use a heater with a PID controller to stabilize the base temperature at 20mK. This mitigates temperature fluctuations during experiments.

We use a 24-port microwave package [1] to host the sample, consisting of a copper casing, a multilayer interposer, a shielding cavity in the package center, and a set of microwave SMP connectors. The signal and ground connections between the package and the sample are provided using superconducting aluminum wirebonds.

Semi-rigid coaxial cables with 0.086" outer conductor diameter and SMA connectorization transmit qubit control and readout signals between the mixing chamber stage and room temperature electronics. Cables connecting stages at different temperatures have low thermal conductivity to minimize the heat flow from warmer stages. We use KEYCOM ULT-05 cables between room temperature and the 4 K stage. All input lines (readout input, qubit control, and TWPA pump lines) use stainless-steel (SS-SS) cables between the 4 K and the mixing chamber stages. Readout output lines use low-loss superconducting niobium-titanium (NbTi) cables between the 4 K and mixing chamber stages to maximize the signal-to-noise ratio of the readout signal. We use hand-formable, non-magnetic, copper coaxial cables (EZ Form Cable, EZ-Flex.86-CU) between the mixing chamber stage and the sample package.

We attenuate and filter the signal at different stages of the dilution refrigerator to reduce the noise incident on the qubits. For the qubit control lines, we use a 20 dB attenuator (XMA Corporation) at the 4 K, a 1 dB attenuator at the still, and a 1 GHz low-pass filter (Mini-Circuits VLF-1000+) at the mixing chamber stage. The 1 dB attenuator at the still was selected to thermalize the center conductor of the coaxial line while maintaining the ability to apply DC current through the line without causing excessive heating. The VLF-1000+ low-pass filter reduces the noise at higher frequencies while allowing an RF drive to reach the qubit. The filter response in the range between 4 – 5 GHz varies by only 3 dB, approximately providing 40 dB of attenuation. Our readout input lines contain a 20 dB attenuator at 4 K, a 10 dB attenuator at the still, and 40 dB of attenuation at the mixing chamber stage, followed by a 12 GHz low-pass filter (RLC F-30-12.4-2). This attenuation and filtering scheme mitigates the thermal-photon population in the resonators [2], which limits the qubit coherence times due to photon shot noise.

The readout circuit in the samples used in our experiment is set up for measuring in reflection mode. The resonator probe tone is coupled to the sample via a circulator (LNF-CIC4.12A). The reflected readout signal gets routed to the readout output chain using the third port of the circulator. The signal is amplified at the mixing chamber stage using a near-quantum-limited Josephson traveling-wave parametric amplifier (TWPA) [3] fabricated at MIT Lincoln Laboratory and mounted inside the μ -metal shield. The TWPA pump tone is combined with the readout signal using a directional coupler (Marki C20-0116). The signal then travels through two isolators (Quinstar 0XE89) followed by a 12 GHz low-pass filter (RLC F-30-12.4-2) and a 3 GHz high-pass filter (RLC F-19704) to prevent noise from coupling back to the sample through the readout chain. Afterward, the signal is once again amplified using a low-noise high-electron-mobility transistor (HEMT) amplifier at the 4 K stage before being transmitted outside the cryostat to room temperature measurement electronics.

We use a microwave source (Rohde & Schwarz SGS100A) with an internal IQ mixer to drive the readout resonators. We use a single-sideband mixing scheme to address multiple resonators coupled to the same feedline using a single microwave local oscillator (LO). The intermediate frequency (IF) signals are generated using two arbitrary waveform generator (AWG) channels (Keysight M3202) and are interfered with the LO using the IQ mixer to obtain an RF signal with multiple frequency components, each of which can address a different resonator. In order to avoid saturating the IQ mixer, we attenuate the output signal of the AWG channels using 10 dB attenuators.

At room temperature, the resonator response signal is amplified by a HEMT amplifier (MITEQ AMF-5D-00101200-23-10P) and sent through a band-pass filter (K&L 033F8) to filter out the TWPA pump tone. The signal is then down-converted using a heterodyne demodulation scheme [4] using an IQ mixer (Marki IQ-4509LXP) and the readout probe LO. The demodulated signals resulting from the interference between the readout output signal and the LO are then filtered using low-pass filters (Mini-circuits VLF-300+), passing only the terms at different frequencies corresponding to the signal from each resonator. The two output signals are then digitized on two separate digitizer channels

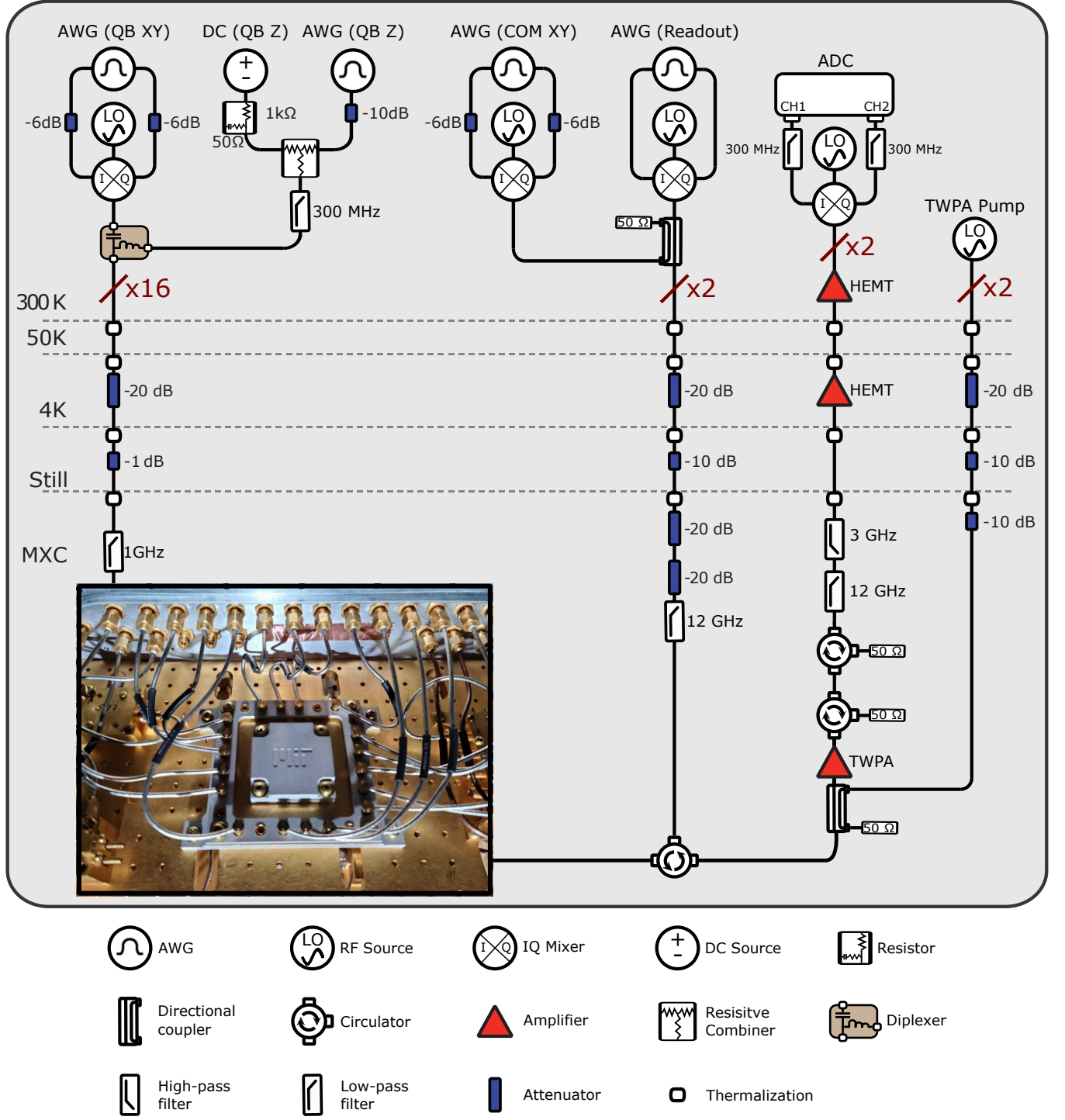


FIG. S1. Measurement setup.

(Keysight M3102), and the signal for each frequency component is processed using a built-in field-programmable gate array (FPGA) to obtain the static in-phase, I , and quadrature, Q , signals.

We use the secondary LO output of the microwave source for the demodulation LO. In order to reduce the signal noise, we use two high-pass filters (Mini-Circuits VHF-6010+) in series. We use a combination of attenuators and an amplifier (Mini-Circuits ZVE-8G+) to ensure that the LO power input to the IQ mixer is in the 10–13 dBm range required for optimal mixer performance.

A 24-channel DC voltage source (QDevil QDAC-I) with in-line resistors generates the current required to apply a static flux to each qubit in our setup. In addition to this static biasing of the qubit, we can apply a baseband, fast-flux pulse to each qubit individually with an AWG channel (Keysight M3202), which allows us to tune the qubit frequency of the qubits on the timescale of nanoseconds. We attenuate the flux pulse produced by the AWG using a 10 dB attenuator in order to reduce the flux noise experienced by the qubit.

For individual X, Y control of each qubit, we use a microwave source (Rohde & Schwarz SGS100A) with an internal IQ mixer. We calibrate the I and Q quadrature offsets, the gain imbalance, and the skew of our IQ mixer by minimizing both the LO leakage and the signal level of the unwanted sideband. We use single-sideband mixing to generate our drive pulses and use an additional AWG channel to gate the RF output in order to mitigate the adverse effects of mixer LO leakage. In our setup, we use four different AWG channels to fully control each qubit: 1 channel for fast Z control, and 3 channels for X, Y control.

We combine the DC bias (used for tuning the static frequency of the qubit), baseband pulse (used for fast tuning of the qubit frequency), and RF pulse (used for charge driving the qubit) at RT. The DC and baseband flux signals are combined using a resistive combiner (Mini-circuits *ZFRSC-42B-S+*). In order to match the impedance of the resistive combiner to $50\,\Omega$ to prevent distortions in the baseband pulse, we use a homemade resistor box on the DC side, consisting of a $1\,\text{k}\Omega$ resistor in parallel with a $50\,\Omega$ resistor in series with a $10\,\mu\text{F}$ capacitor connected to ground. The output flux signal is sent through a 300 MHz low-pass filter (Mini-circuit VLFX-300+) and is combined with the RF signal using a diplexer (QMC-CRYODPLX-0218). We use a 1 GHz low-pass filter (VLFG-1000+) with a flat response of approximately 40 dB of rejection throughout the qubit operation frequency range. Using this filter, we are able to filter out high-frequency sources of noise while maintaining the ability to apply high-fidelity single-qubit gates.

The advantage of this approach is that we only need a single microwave coaxial line for full control of each qubit. In addition, by combining DC and baseband flux signals at RT we can more efficiently characterize and compensate for the flux transients arising from signal combination with a fast oscilloscope or a digitizer. However, the impedance mismatch caused by the reflective low-pass filter at the mixing chamber gives rise to standing modes between the filter and the chip ground at the end of the flux line. As a result, the effective qubit-drive coupling depends on the drive frequency. We characterize this dependence by biasing the qubit at different frequencies and Rabi driving on resonance.

The active components used in our experimental setup are summarized in Table SI.

Quantity	Electrical component	Part number
2×	PXI chassis	Keysight M9019
18×	AWG	Keysight M3202
2×	Digitizer	Keysight M3102
18×	Microwave source	Rohde & Schwarz SGS100A
1×	Microwave source	Holworth HS9004A (4 channels)
1×	DC source	QDevil QDAC-I (24 channels)
1×	Reference clock	SRS Model FS725
1×	Digitizer pre-amplifier	SR445A 350MHz pre-amplifier
2×	Amplifier	MITEQ AMF-5D-00101200-23-10P
1×	Demodulation LO amplifier	Mini-Circuits ZVE-8G+
2×	TWPA	Made at MIT LL

TABLE SI. **Active electrical components used in the experimental setup.**

S2. SAMPLE

The 4×4 superconducting transmon array is fabricated using a 3D integration flip-chip process [5] (Fig. 1f in the main text). The qubits are located on a 4×4 mm qubit tier as shown in Fig. 1g, and the readout and control lines are located on a separate 5×5 mm interposer tier as shown in Fig. 1h. The two layers are separated using $3\,\mu\text{m}$ superconducting indium bump bonds and silicon hard stops.

The interposer tier contains the flux bias lines and readout resonators that are respectively inductively and capacitively coupled to the qubits across the gap separating the two chips. The ground plane of the interposer layer is etched away in the regions corresponding to the location of each qubit, with an additional $10\,\mu\text{m}$ gap to avoid unwanted capacitance between the qubits and the interposer ground plane.

The qubit tier comprises a 2D lattice of capacitively-coupled single-ended transmon qubits [6] arranged in 4×4 array. The capacitive coupling between each transmon is facilitated using an intermediary piece of superconductor, resulting

in a nearest-neighbor qubit-qubit exchange interaction with average strength $J/2\pi = (5.89 \pm 0.4)$ MHz, measured at qubit frequencies of 4.5 GHz. The ground plane of the qubit tier is etched away in the regions corresponding to the location of the signal line on the interposer tier, with an additional gap equal to the width of each line.

Multiplexed, dispersive qubit-state readout is performed through individual capacitively coupled coplanar resonators. The sample contains two readout feedlines, each containing a single-port Purcell filter coupled to eight readout resonators. We use $\lambda/4$ readout resonators for all qubits along the edge of the lattice to minimize their footprint on the sample. The central qubits are read out with a $\lambda/2$ resonator. This is done to minimize the parasitic couplings; the central qubits' resonator crossing is located in the middle of the $\lambda/2$ resonator at its voltage node. The Purcell filter is designed to have a large bandwidth of 0.5 GHz at a resonance frequency of 6.3 GHz, such that the readout resonators can be distributed over a ~ 400 MHz band while limiting Purcell decay to a rate $\lesssim 1/(300 \mu\text{s})$. The resonators have a measured average linewidth of $\kappa/2\pi = (1.17 \pm 0.3)$ MHz, with an average dispersive shift of $\chi/2\pi = (0.87 \pm 0.11)$ MHz characterized with each qubit biased at 4.5 GHz.

The 3D integration process allows us to route each flux bias line directly above the corresponding superconducting quantum interference device (SQUID). This positioning allows us to attain an average mutual inductance between the qubit SQUID loops and their respective flux bias lines of 1.15 pH while reducing the SQUID area to $4 \times 4 \mu\text{m}$ compared to planar transmon arrays [7, 8]. Decreasing the area of the SQUID reduces the susceptibility of the qubit to flux noise, the offset due to uncontrollable magnetic fields in the environment, and crosstalk. We measure more than an order of magnitude reduction in the crosstalk levels of the flip-chip sample compared to a planar 3×3 array of transmons [7, 8].

The sample is fabricated on a silicon substrate by dry etching an MBE-grown, 250-nm-thick aluminum film in an optical lithography process, forming all larger circuit elements such as the qubit capacitor pads, resonators, and the signal lines for qubit readout and control. The qubit SQUID loops are fabricated with an electron beam lithography process and a double-angle shadow evaporation technique [9] to form the Josephson junctions. We use a wire width of $5 \mu\text{m}$ for the SQUIDs to minimize flux noise from local magnetic spin defects on the SQUID surface and interfaces [10].

Detailed sample parameters for individually isolated qubits are summarized in Table S2. We show the maximum transmon transition frequencies ω_q^{max} at the upper flux insensitive point, the qubit anharmonicities U (measured at ω_q^{max}), the readout resonator frequencies ω_{res} , the probabilities $\mathcal{F}_{ij}$ of measuring the qubit in state i after preparing it in state j , the readout assignment fidelity $(\mathcal{F}_{\text{gg}} + \mathcal{F}_{\text{ee}})/2$, the measured T_1 s and T_2^* s at qubit frequencies around the bias point used in the experiment. We anticipate that the reported readout assignment fidelities are limited by the state preparation error due to the initial thermal population of the qubits. The relatively low $\mathcal{F}_{\text{gg}}$ is a result of the high thermal population caused by noticeable heating at the MXC stage when DC biasing the qubits. This heating is caused by sending a DC current through the microwave cable between the still and MXC stages, dissipating approximately $\sim 500 \mu\text{W}$ of power at the MXC when all the qubits are biased at 4.5 GHz. In addition, we report the individual and simultaneous single-qubit randomized benchmarking (RB) fidelity for the gates used for tomographic readout. The microwave charge pulses for the gates are applied via the flux lines. Due to the significant RF crosstalk between the drives targeting QB5 and QB10, we do not calibrate and apply simultaneous gates for tomographic readout on these qubits.

S3. CONTROL AND READOUT CALIBRATION

A. Flux crosstalk calibration

Flux-tunable transmon qubits comprise two Josephson junctions in a SQUID loop in parallel with a shunting capacitor. We can control the effective Josephson energy E_J of the SQUID, and therefore the transition frequency between the ground and the first-excited state of the transmon, by threading magnetic flux through the SQUID. The frequency of the transmon in response to an applied external flux Φ_{ext} is approximately given by [6]:

$$f(\Phi_{\text{ext}}) = \left(f^{\text{max}} + \frac{E_C}{h} \right) \sqrt[4]{d^2 + (1 - d^2) \cos^2 \left(\pi \frac{\Phi_{\text{ext}}}{\Phi_0} \right)} - \frac{E_C}{h}, \quad (1)$$

where E_C is the transmon charging energy and $f^{\text{max}} = (\sqrt{8E_J E_C} - E_C)/h$ is the maximum qubit frequency. The asymmetry parameter d of the SQUID junctions is given by $d = |(E_{J,2} - E_{J,1})/(E_{J,2} + E_{J,1})|$, where $E_{J,1}$ and $E_{J,2}$ are the Josephson energies of the two SQUID junctions. The above equation is an approximation because it holds only in the limit $E_J \gg E_C$.

We thread flux through the SQUID by applying a current through a local flux line that terminates near the SQUID. We use a voltage source applied over a series resistor of resistance $1 \text{ k}\Omega$ at room temperature to generate a stiff current

Parameters	QB1	QB2	QB3	QB4	QB5	QB6	QB9	QB10
$\omega_q^{\max}/2\pi$ (GHz)	4.867	4.89	4.801	4.845	4.706	4.833	5.061	4.972
$U/2\pi$ (MHz)	-217.3	-217.6	-217.0	-241.9	-217.9	-214.5	-217.9	-210.1
$\omega_{\text{res}}/2\pi$ (GHz)	6.207	6.358	6.237	6.496	6.367	6.337	6.285	6.431
$\mathcal{F}_{\text{gg}}$	0.96	0.96	0.96	0.95	0.96	0.95	0.95	0.96
$\mathcal{F}_{\text{ee}}$	0.92	0.88	0.92	0.93	0.89	0.88	0.86	0.81
Readout assignment fidelity	0.94	0.93	0.94	0.94	0.94	0.91	0.91	0.9
T_1 (μs) at operating point	23.8	24.1	16.7	34.9	17.6	15.2	14.7	17.1
T_2^* (μs) at operating point	3.6	2.7	1.8	2.6	2.9	3.8	2.7	1.1
Single-qubit RB (individual)	99.78%	99.88%	99.84%	99.89%	NA	99.84%	99.88%	NA
Single-qubit RB (simultaneous)	99.42%	99.88%	99.72%	99.71%	NA	99.81%	99.87%	NA

Parameters	QB7	QB8	QB11	QB12	QB13	QB14	QB15	QB16
$\omega_q^{\max}/2\pi$ (GHz)	4.699	4.944	4.892	4.865	5.077	5.009	4.785	4.948
$U/2\pi$ (MHz)	-217.8	-217.6	-213.3	-219.2	-218.6	-216.6	-219.3	-218.4
$\omega_{\text{res}}/2\pi$ (GHz)	6.425	6.28	6.336	6.356	6.503	6.25	6.359	6.196
$\mathcal{F}_{\text{gg}}$	0.95	0.95	0.96	0.95	0.97	0.96	0.96	0.96
$\mathcal{F}_{\text{ee}}$	0.92	0.91	0.9	0.91	0.83	0.9	0.9	0.87
Readout assignment fidelity	0.93	0.93	0.93	0.94	0.9	0.94	0.93	0.93
T_1 (μs) at operating point	22.1	25.6	26	23.5	22.1	25.9	28.2	26.7
T_2^* (μs) at operating point	2.8	2.7	2.3	2.1	2.4	1.3	2.4	1.1
Single-qubit RB (individual)	99.78%	99.82%	99.88%	99.86%	99.77%	99.84%	99.9%	99.86%
Single-qubit RB (simultaneous)	99.63%	99.73%	99.78%	99.14%	99.83%	99.57%	99.39%	99.34%

TABLE SII. **Sample parameters.** We show the maximum transmon transition frequencies $\omega_q^{\max}$ at the upper flux-insensitive point, the qubit anharmonicities U (measured at $\omega_q^{\max}$), the readout resonator frequencies ω_{res} , the probabilities $\mathcal{F}_{ij}$ of measuring the qubit in state i after preparing it in state j , the readout assignment fidelity $(\mathcal{F}_{\text{gg}} + \mathcal{F}_{\text{ee}})/2$, the measured T_1 s, T_2^* s. The reported readout assignment fidelities are limited by the state preparation error due to the initial thermal population of the qubits. In addition, we report the individual and simultaneous single-qubit randomized benchmarking (RB) fidelities. Due to the significant RF crosstalk between the drive targeting QB5 and QB10, we do not apply single-qubit gates to these qubits in our experiments.

source for each flux line. There is a linear relation between the voltage applied to a flux line and the magnetic flux experienced by the SQUID:

$$\Phi_{\text{ext}} = V/V^{\Phi_0} + \Phi_{\text{offset}}, \quad (2)$$

where V^{Φ_0} is the voltage required to tune the qubit by one magnetic flux quantum Φ_0 , and Φ_{offset} is a flux offset due to non-controllable sources of static magnetic field.

For an array of transmons, the applied current in a flux line targeting one SQUID may thread unwanted magnetic flux through other SQUIDS. We can model this flux crosstalk as a linear process, where the fluxes $\vec{\Phi}_{\text{ext}}$ experienced by the SQUIDS are related to the voltages $\vec{V}$ applied to the flux lines:

$$\vec{\Phi}_{\text{ext}} = (\mathbf{V}^{\Phi_0})^{-1} \mathbf{S} \vec{V} + \vec{\Phi}_{\text{offset}}, \quad (3)$$

where $\mathbf{V}^{\Phi_0}$ is a diagonal matrix with $V_{i,i}^{\Phi_0}$ corresponding to the V^{Φ_0} of qubit i . $\mathbf{S}$ is the flux crosstalk matrix, with $S_{i,j} = \partial V_i / \partial V_j$ representing the voltage response of qubit i to a voltage signal applied to qubit j . In this

representation, the diagonal elements of $\mathbf{S}$ are 1. By characterizing $\mathbf{S}$, we can compensate for the crosstalk and set the qubit frequencies precisely.

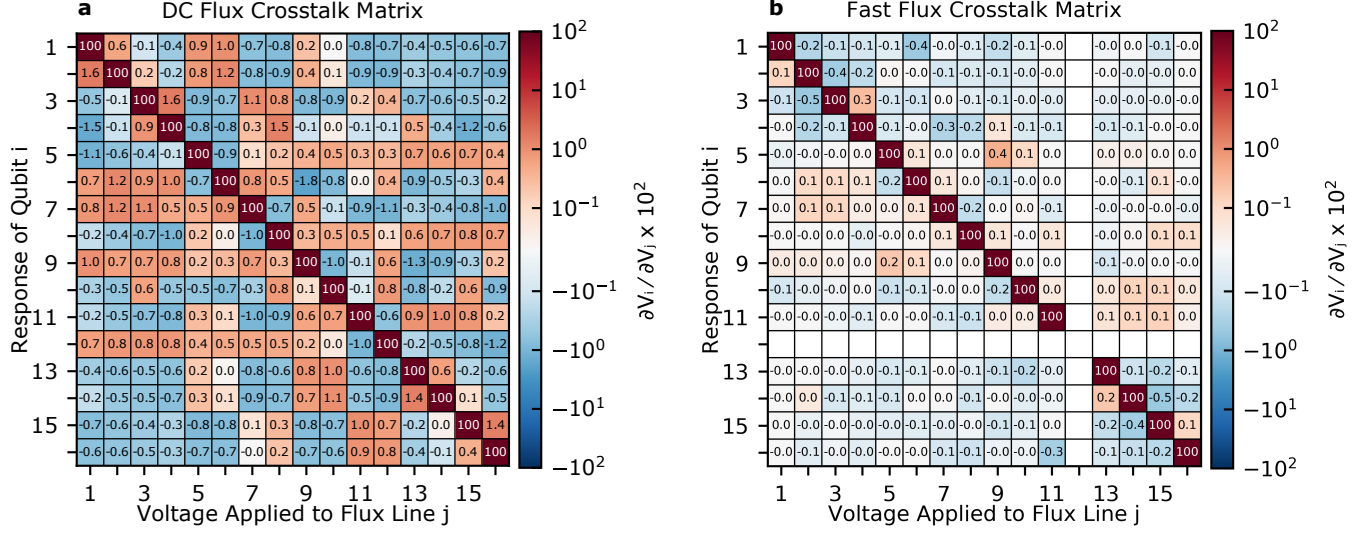


FIG. S2. **Learned Flux Crosstalk Matrices.** The learned (a) DC flux crosstalk matrix and (b) fast flux crosstalk matrix for the 16-qubit device, each rescaled by a factor of 100 such that each element is a percentage. The fast flux microwave line for qubit 12 was broken, so we have no crosstalk information for qubit 12.

The first step of calibrating flux crosstalk is characterizing the transmon spectrum of each qubit in the array. We do this by sweeping the voltage applied to the flux line and measuring the qubit frequency. By fitting this data with Eq. 1 and Eq. 2, we extract the transmon spectrum parameters $f^{\max}$, E_C , and d , and the voltage to flux conversion parameters V/V^{Φ_0} and Φ_{offset} .

Next, we calibrate flux crosstalk with a learning-based protocol described in [11]. We start with an initial guess crosstalk matrix, which can either be an estimate from a previous measurement or the identity matrix.

We select a quasi-random vector of target frequencies $\vec{f}$, subject to a few constraints. First, we require that the frequency for each qubit falls in the range of 100 MHz below its sweet spot to 1 GHz below its sweet spot. Second, to reduce frequency shifts due to resonant exchange interaction between qubits, we require that the frequency detuning between neighboring qubits is at least 200 MHz and that the frequency detuning between any two qubits in the array is at least 50 MHz. We use our initial $\mathbf{S}$ and equations Eq. 1 and Eq. 3 to solve for the voltages $\vec{V}$ needed to target $\vec{f}$. We apply $\vec{V}$, then measure the qubit frequencies via spectroscopy. Even though we detune our qubits, they still experience frequency shifts. We use previously characterized qubit-qubit couplings to calculate the uncoupled qubit frequencies. Finally, we convert the uncoupled frequencies to fluxes experienced.

By repeating this process M times, we obtain a size- M training set of applied voltages and measured fluxes: $\{\vec{V}_i, \vec{\Phi}_{\text{meas},i}\}_{i=1:M}$. We learn our crosstalk matrix using the crosstalk relation (Eq. 3), by leveraging the fact that the estimated flux $(\mathbf{V}^{\Phi_0})^{-1} \mathbf{S} \vec{V}_i + \vec{\Phi}_{\text{offset}}$ ought to equal the measured flux $\vec{\Phi}_{\text{meas},i}$. We learn $\mathbf{S}$ row-by-row. We learn the elements of the k^{th} row of $\mathbf{S}$ by minimizing the mean-squared-error cost function

$$C(\mathbf{S}_k) = \frac{1}{M} \sum_{i=1}^M \left\| (\vec{\Phi}_{\text{meas},i})_k - \left[(\mathbf{V}_{k,k}^{\Phi_0})^{-1} \mathbf{S}_k \vec{V}_i + (\vec{\Phi}_{\text{offset}})_k \right] \right\|^2. \quad (4)$$

in a gradient descent optimizer (we use the L-BFGS optimizer in PyTorch). The optimization will converge as the estimated fluxes approach the measured fluxes. Once we have repeated this minimization for all rows, we have our optimized $\mathbf{S}$. In Fig. S2a, we report the DC flux crosstalk matrix for the 16-qubit device.

This same calibration procedure is utilized to calibrate crosstalk for fast flux pulse control, as well. In order to learn the fast flux crosstalk for each qubit, we tune that qubit to a target frequency using a 100 ns fast flux pulse and measure the qubit frequency via spectroscopy. Then, using the same pulse amplitude for the target qubit, we apply flux pulses with random amplitudes through the other flux lines and measure the change in the target qubit's frequency. We use the training set consisting of the voltages applied and the changes in the target qubit frequency for learning the fast flux crosstalk matrix. In Fig. S2b, we report the fast flux crosstalk matrix for the 16-qubit device.

B. Time-domain flux pulse shaping and transient calibration

Baseband-frequency flux pulses experience distortions as they travel to the sample. The distortion caused by room-temperature components can be characterized and compensated for using a fast oscilloscope with a sampling rate better than 500 MHz. The distortion introduced by the components in the cryostat is generally temperature-dependent and needs to be characterized while the fridge is operating. By using the qubit as a sensor, we can characterize the distortion of a square flux pulse (Z pulse) using Ramsey-type measurements [12]. The step response for the signal generated by the AWG, $s_{\text{AWG}}(t)$, the signal reaching the qubit, $s_{\text{qubit}}(t)$, and the time-dependent distortion, $n(t)$, are related by

$$s_{\text{qubit}}(t) = s_{\text{AWG}}(t) (1 + n(t)), \quad (5)$$

To characterize $n(t)$, the target qubit is biased to its sweet spot using static flux and excited with a $\frac{\pi}{2}$ -pulse around the x -axis followed by a Z pulse for time t . The dynamic frequency change of the qubit as a response to the Z pulse can be extracted by measuring both $\langle X \rangle$ and $\langle Y \rangle$, denoting the expectation values of the qubit state projected along the x -axis and y -axis of the Bloch sphere. The phase accumulated by the qubit during the Z pulse can be described by $\phi(t) = \int_0^t \Delta(t') dt'$, where $\Delta(t) = \Delta_0 + \delta(t)$ is the frequency detuning, consisting of the target detuning Δ_0 and the added frequency distortion $\delta(t)$. The $\langle X \rangle$ and $\langle Y \rangle$ measurements are used to obtain $\cos \phi(t)$ and $\sin \phi(t)$, respectively, which we use for constructing the phasor $e^{i\phi(t)} = \cos \phi(t) + i \sin \phi(t)$. We can then extract the frequency distortion as $\delta(t) = \frac{d}{dt} \arg(e^{i\phi(t)})$. By using the qubit spectrum, we map $\delta(t)$ to the distortion in the signal reaching the qubit $n(t)$.

In order to obtain an analytical expression, we fit the extracted $n(t)$ with a sum of typically three or more exponential damping terms with time constants τ_i and settling amplitudes A_i :

$$n(t) = \prod_i A_i e^{-t/\tau_i} \quad (6)$$

In our experiments, we use the analytical expression of $n(t)$ to pre-distort the output signal of the AWG in order to compensate for the system transients.

One method of analytically pre-distorting the signal is using Fourier transformations. Using the signal step-response $h(t) = \Theta(t)(1 + n(t))$, we can obtain the impulse-response of our system:

$$I(t) = \dot{h}(t) = \delta(t) + \delta(t) \prod_i A_i e^{-t/\tau_i} - \Theta(t) \sum_j \left(\frac{1}{\tau_j} \prod_i A_i e^{-t/\tau_i} \right).$$

We can solve for the pre-distorted signal needed to make the qubit feel a constant flux using the convolution relation $s_{\text{qubit}}(t) = s_{\text{AWG}}(t) * I(t)$, therefore,

$$s_{\text{AWG}}(t) = \mathcal{F}^{-1} \left\{ \frac{\mathcal{F}\{s_{\text{qubit}}(t)\}}{\mathcal{F}\{I(t)\}} \right\}$$

where $\mathcal{F}\{I(t)\} = \mathcal{F}\{\delta(t)\} + \mathcal{F}\{\delta(t) \prod_i A_i e^{-t/\tau_i}\} - \sum_j \frac{1}{\tau_j} \mathcal{F}\{\Theta(t) \prod_i A_i e^{-t/\tau_i}\}$. We can calculate an expression for $\mathcal{F}\{I(t)\}$:

$$\mathcal{F}\{I(t)\} = \prod_i \left(1 + \frac{i\omega A_i \tau_i}{i\omega \tau_i + 1} \right).$$

In practice, in order to accurately pre-distort the pulse shape using Fourier transformations we need to zero-pad the waveform such that the total length of the waveform is at least $5 \times \max_i(\tau_i)$. As a result, when correcting for long-timescale transients with $\tau \propto \mathcal{O}(100 \mu\text{s})$ this procedure becomes inefficient.

Instead, we use an infinite-impulse-response (IIR) filter for pulse pre-distortion [12] with the difference equation:

$$a_0 y[n] = \sum_{i=0}^N b_i x[n-i] - \sum_{j=1}^M a_j y[n-j] \quad (7)$$

where $y[n]$ is the output signal, $x[n]$ is the input signal, N is the feed-forward filter order, and M is the feed-back filter order. A first-order IIR filter ($N = M = 1$) can be used to correct the exponential transients. We find the IIR coefficients a_0 , a_1 , b_0 , and b_1 for a single exponential pole with amplitude A and time-constant τ :

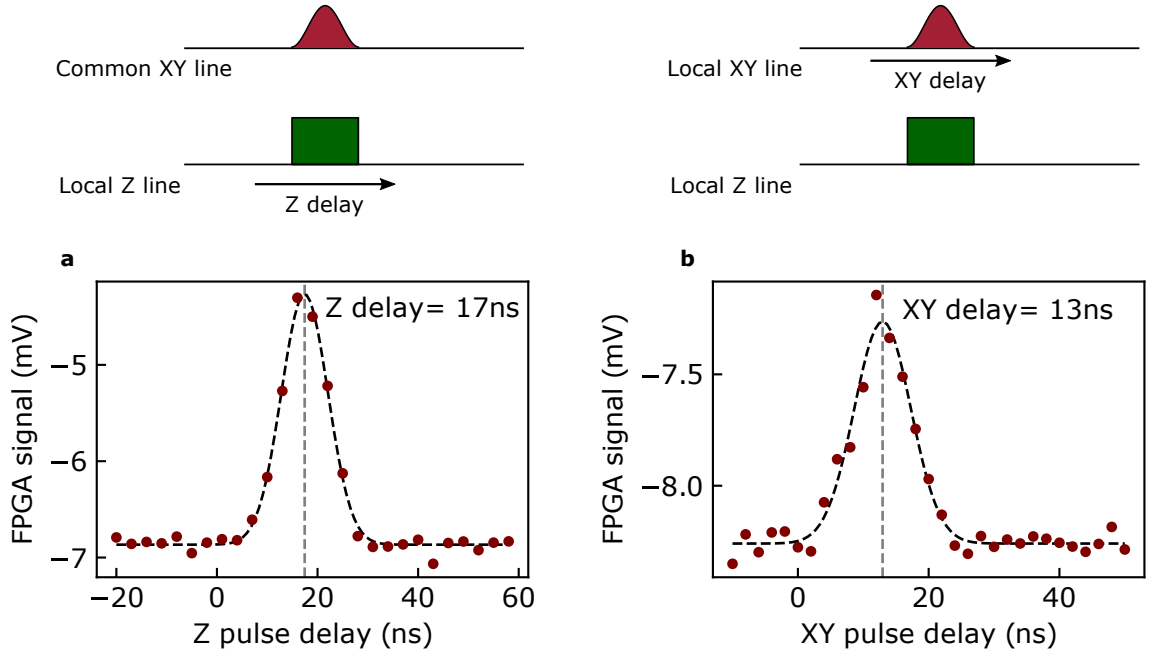


FIG. S3. **Pulse alignment calibrations.** (a) The timing for the Z pulse applied through each line is calibrated by sweeping the delay time of a Z pulse by tuning the qubit to the sweet spot while applying an XY pulse through the common resonator feedline at that frequency. (b) The timing for the XY pulse applied through each local line is similarly calibrated by applying a Z pulse with a constant delay characterized in (a) and sweeping the delay time of the XY pulse.

$$a_0 = 1 + A, \quad a_1 = -\alpha - A, \quad b_0 = 1, \quad b_1 = -\alpha, \quad (8)$$

where $\alpha = e^{-1/f_s \tau}$ depends on the AWG sampling rate f_s .

We use the SciPy library in Python to apply the IIR filter with the coefficient calculated above to the waveforms in our pulse generation software. For systems with more than one exponential transient term, we concatenate multiple IIR filters, each corresponding to a term.

C. Pulse alignment

Differences in the pulse path lengths and the delays in control electronics cause discrepancies in the time that flux (Z) and charge (XY) pulses reach the target qubit. To accurately execute the desired pulse sequence, we characterize the relative delay of each signal line and adjust pulse timings in software to compensate.

We first characterize the delay of a Z pulse applied to each qubit with respect to a common XY drive through the resonator feedline. An XY pulse through the resonator feedline reaches each qubit roughly simultaneously; it, therefore, serves as a common reference for pulse delay timing. To calibrate the Z pulse delay of a particular qubit, we first isolate the qubit from the rest of the lattice and away from the flux sweet spot. We then apply a short 13 ns Z pulse with the amplitude required to tune the qubit to the sweet spot, and an XY pulse of the same length at the sweet-spot frequency. The amplitude of the XY pulse is less than that required for a π -pulse. We measure the qubit resonator response as a function of the Z pulse delay (Fig. S3a); a peak in the measured response indicates the delay at which the qubit is maximally excited, which occurs when the two pulses are aligned. (When the two pulses are misaligned, part, or all, of the XY drive is off-resonant with the qubit frequency, resulting in a lower qubit excited state population.) By fitting the data in Fig. S3a to a Gaussian curve, we find the optimal Z delay of 17 ns with respect to the XY pulse applied through the common line.

Next, we calibrate the delay for XY pulses sent to each qubit through its respective local line, using the calibrated Z pulse delays for reference. Similar to the procedure described above, we tune the qubit to the sweet spot using a Z -pulse and apply an XY pulse, now through the local line, at the qubit sweet spot frequency. Now with the Z pulse timing fixed at the calibrated value, we sweep the delay of the local XY pulse and extract the optimal local XY pulse delay with respect to the common XY pulse timing, as shown in Fig. S3b.

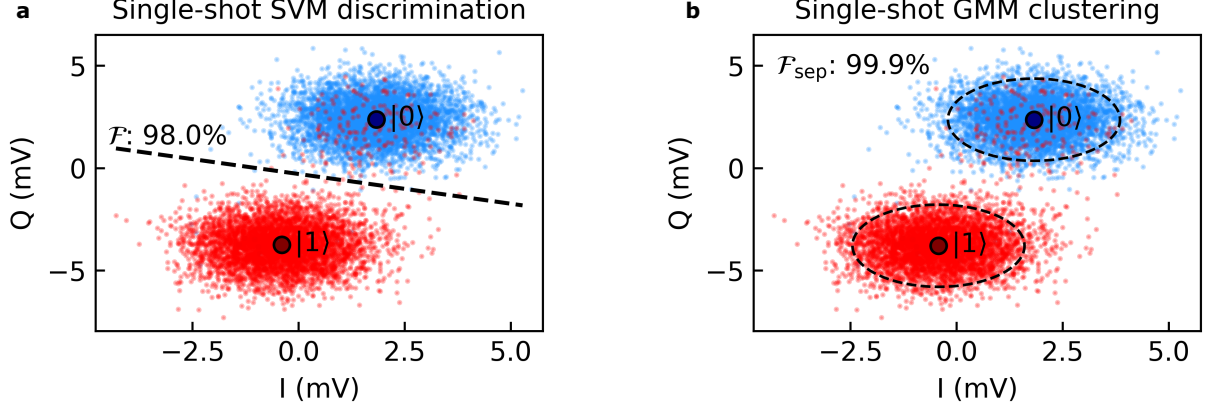


FIG. S4. **Readout single-shot data.** (a) We use an SVM to draw a discrimination line (dashed line) between the single-shot points prepared in the $|0\rangle$ state (blue points) and $|1\rangle$ state (red points) prior to readout. From this data, we obtain the readout fidelity $\mathcal{F}_{\text{RO}} = 98.0\%$. (b) The same single-shot dataset is grouped into two clusters using a GMM with a cluster separation fidelity of $\mathcal{F}_{\text{sep}} = 99.9\%$

D. Readout optimization

We use frequency multiplexed heterodyne mixing to read out multiple resonators coupled to the same feedline with a common local oscillator (LO) frequency [4]. To maximize the readout fidelity, we maximize the separation in in-phase/quadature (I/Q) space between the readout signal conditioned on the qubit ground and excited state by optimizing the frequency and amplitude of each sideband tone.

In order to distinguish between the single-shot clusters resulting from different qubit states we first take single-shot measurements with the qubit prepared in the $|0\rangle$ state and the $|1\rangle$ state by applying a π -pulse to the qubit. Using this data, we train a support vector machine (SVM) [13] to find the hyperplane that classifies each single-shot point as 0 or 1. The hyperplane parameters are calculated in a manner that maximizes the probability of correctly classifying each data point. We quantify the readout quality using the readout fidelity:

$$\mathcal{F}_{\text{RO}} = \frac{f_{00} + f_{11}}{2}, \quad (9)$$

where $f_{00} = P(0| |0\rangle)$ is the probability of correctly classifying the qubit $|0\rangle$ state as 0, and $f_{11} = P(1| |1\rangle)$ is the probability of classifying the $|1\rangle$ state as 1 after readout.

Fig. S4a illustrates an example set of single-shot data acquired by preparing the qubit in the $|0\rangle$ state (blue) and in the $|1\rangle$ state (red) and the corresponding line used to discriminate between the two states. For this dataset, the readout fidelity is $\mathcal{F}_{\text{RO}} = 98.0\%$, with $f_{00} = 99.6\%$ and $f_{11} = 96.5\%$. The discrepancy between f_{11} and f_{00} is a result of qubit relaxation during the readout process.

In addition to using $\mathcal{F}_{\text{RO}}$ to quantify the readout quality, we also consider cluster separation. We use a Gaussian Mixture Model (GMM) probability distribution to group the single-shot data into two clusters, independent of the intended state. We fit the mean (μ) and the covariance matrix Σ for each Gaussian distribution representing a cluster using the expectation-maximization (EM) algorithm [13]. A multivariate Gaussian distribution in the I/Q plane with mean μ and covariance matrix Σ has density:

$$P_{\mu, \Sigma}(\mathbf{x}) = \frac{1}{2\pi\sqrt{|\Sigma|}} \exp\left(-\frac{1}{2}(\mathbf{x} - \mu)^T \Sigma^{-1}(\mathbf{x} - \mu)\right). \quad (10)$$

In our model, we use the same covariance matrix Σ for both clusters. We use the spatial overlap between the Gaussian distributions to quantify the cluster separation by defining the separation fidelity $\mathcal{F}_{\text{sep}}$ as

$$\mathcal{F}_{\text{sep}} = 1 - \int \min[P_{\mu_0, \Sigma}(\mathbf{x}), P_{\mu_1, \Sigma}(\mathbf{x})] d\mathbf{x}. \quad (11)$$

In Fig. S4b, we group the readout single-shot readout into two clusters. The ellipse around each cluster encloses the area corresponding to two standard deviations away from the center of the cluster. Using Eq. 11 we obtain a separation fidelity of $\mathcal{F}_{\text{sep}} = 99.9\%$ between the two readout clusters.

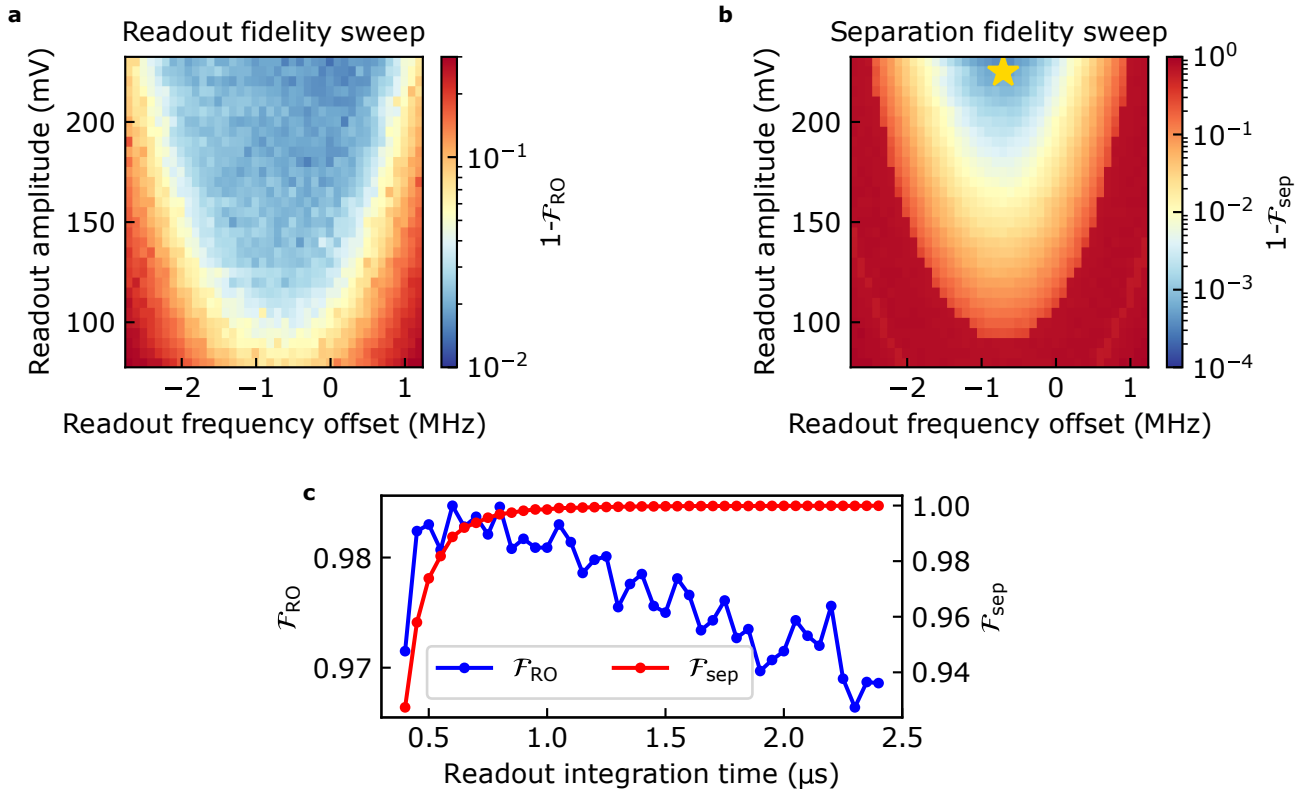


FIG. S5. **Readout optimization.** (a) Experimentally measured readout optimization landscape using the readout fidelity. The plateau in the quantity is caused by qubit relaxation during readout. (b) Experimentally measured readout optimization landscape using the cluster separation fidelity. The star depicts the optimal parameters. (c) $\mathcal{F}_{\text{RO}}$ and $\mathcal{F}_{\text{sep}}$ at different readout integration times.

To optimize the readout in our calibration procedure, we vary the sideband modulation pulse frequency and amplitude used for multiplexed readout. Both the readout fidelity $\mathcal{F}_{\text{RO}}$ and the separation fidelity $\mathcal{F}_{\text{sep}}$ can be used as the optimization metric. We measure the readout optimization landscape by sweeping the readout frequency offset and readout amplitude for a given qubit and calculating $1 - \mathcal{F}_{\text{RO}}$ (Fig. S5a) and $1 - \mathcal{F}_{\text{sep}}$ (Fig. S5b). The optimization landscape using both metrics is convex, which allows for the use of a conventional optimizer, such as Nelder-Mead, to find the optimal readout sideband frequency and amplitude efficiently. The landscape of the cluster separation fidelity is more sensitive to changes in the readout parameters compared to the readout fidelity, which is a result of the limit imposed on the readout fidelity caused by qubit relaxation during readout. Therefore, separation fidelity is a better metric to use for the optimal readout sideband frequency and amplitude while keeping the readout time fixed. It is important to limit the maximum readout amplitude to ensure that the readout remains in the dispersive regime. In the dispersive readout scheme, with the readout resonator frequency above the transmon frequency, we find the magnitude of the optimal readout frequency offset to be roughly equal to the measured qubit-resonator dispersive shift χ as predicted for $\chi \lesssim \kappa/2$ [4].

Using these optimal parameters, we vary the readout integration time to maximize the readout fidelity (Fig. S5c). Increasing the integration time results in an increase in the single-shot cluster separation, however, the readout fidelity suffers more from qubit relaxation. Therefore, the optimal integration time is selected in a way to balance these two effects by maximizing $\mathcal{F}_{\text{RO}}$.

E. State preparation and measurement (SPAM) error mitigation

In order to mitigate SPAM errors from our measured observable, we use a stochastic β -matrix (often also referred to as a T matrix in the literature) for each qubit, where

$$\beta = \begin{pmatrix} f_{00} & 1 - f_{11} \\ 1 - f_{00} & f_{11} \end{pmatrix} \quad (12)$$

where f_{00} and f_{11} are the probabilities of correctly preparing and measuring the qubit in the $|0\rangle$ and $|1\rangle$ states respectively. There are two main sources of our experimental error captured by this matrix: incoherent thermal population of the qubit prior to the experiment sequence, and readout errors.

In the absence of readout crosstalk, the β -matrix for an n -qubit system is the tensor product of the β -matrix for each individual qubit: $\bar{\beta} = \beta_1 \otimes \dots \otimes \beta_n$. By inverting this matrix, and applying it to the measured bitstring probabilities we can correct for SPAM errors in the measured observable [14].

F. Idling frequency layout

During the emulation of the Bose-Hubbard Hamiltonian, all qubits are brought on resonance to allow particle exchange. During state preparation and readout steps of the experiment sequence, the qubits are detuned from one another to turn off particle exchange. The capability to apply simultaneous high-fidelity single-qubit gates to qubits at their respective idling frequencies is necessary for performing tomographic readout, which is required to measure correlations and state purities. The idling frequencies of the qubits during the state preparation and readout steps of the experimental sequence satisfy the following criteria:

- Each qubit is detuned from its nearest neighbors by at least 300 MHz to prevent particle exchange during readout.
- Each qubit is detuned from its next-nearest neighbors by at least 40 MHz.
- Qubit frequencies are selected to minimize charge-driving crosstalk effects to enable high-fidelity simultaneous single-qubit gates.
- Qubit 12 is biased at 4.5 GHz (qubit 12 suffered from elevated pulse distortion due to a partial break in its flux line).
- Each qubit frequency is detuned from any two-level system (TLS) defects.

The idling qubit frequency layout used in our experiment is shown in Fig. S6. At these frequencies, we apply single-qubit gates on all qubits, excluding qubits 5 and 10, for which the qubit was too weakly coupled to the line used for charge driving. The randomized benchmarking (RB) fidelities for the calibrated gates are reported in Tab. S2, with an average individual RB fidelity of 99.8% and an average simultaneous RB fidelity of 99.6%.

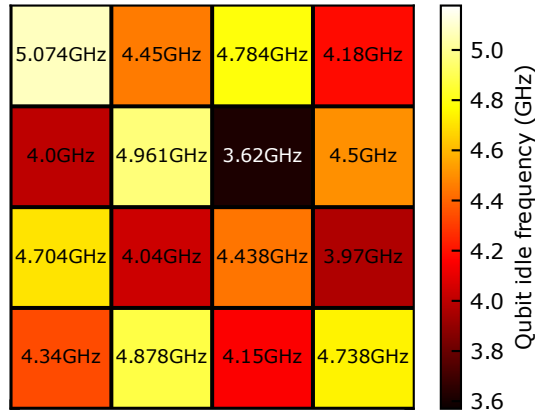


FIG. S6. Lattice qubit idle frequencies.

S4. HAMILTONIAN CHARACTERIZATION

The Hamiltonian of the driven system is described by

$$\hat{H}/\hbar = \sum_{\langle i,j \rangle} J_{ij} \hat{\sigma}_i^+ \hat{\sigma}_j^- + \frac{\delta}{2} \sum_i \hat{\sigma}_i^z + \Omega \sum_j (\alpha_j \hat{\sigma}_j^- + \text{h.c.}), \quad (13)$$

as shown in Eq. 5 of the main text. The particle exchange strengths J_{ij} and the amplitudes and phases of the relative drive couplings α_i are determined by the circuit layout, and remain constant in our experiments. To accurately simulate the behavior of our lattice, we carefully characterize each of these parameters.

A. Particle exchange strengths

We measure the particle exchange strengths J_{ij} for nearest and next-nearest neighboring qubits in the lattice (that is, for qubits i and j of relative Manhattan distance 1 and 2).

Each value J_{ij} is measured via the resonant iSWAP rate between qubits i and j at $\omega_{\text{com}}/2\pi = 4.5$ GHz. With all other qubits detuned, we first add an excitation to qubit i by applying a π -pulse. Qubits i and j are then brought on resonance at frequency ω_{com} . After allowing the qubits to interact for time t , the qubits are once again detuned and measured in the z basis. The coupling strength J_{ij} is determined by the time T required for a full excitation swap, $J = \pi/2T$ [15].

A histogram of the measured values J_{ij} for (next-)nearest neighbors are shown in Fig. S7(a(b)). The average nearest coupling in our device is $J/2\pi = 5.9(4)$ MHz at $\omega_{\text{com}}/2\pi = 4.5$ GHz. The next-nearest neighbor couplings are typically in the range of $J/10$ to $J/40$. The qubits at the center of the lattice exhibit a larger coupling to their next-nearest neighbor, which can be attributed to larger stray capacitances.

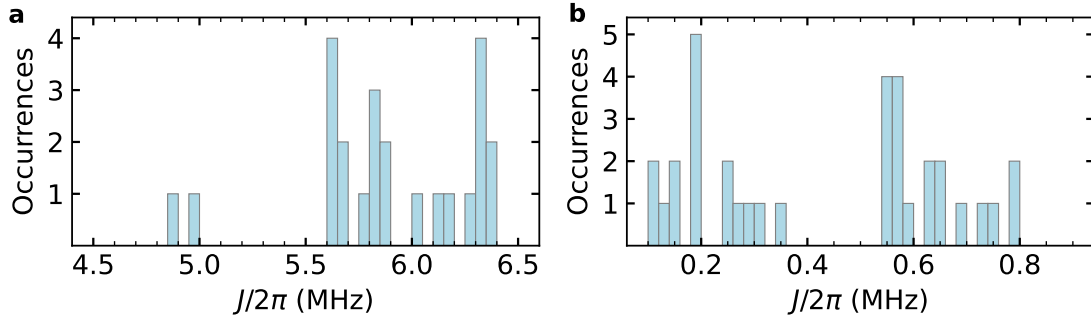


FIG. S7. Histogram of coupling strengths among (a) nearest-neighboring sites and (b) next-nearest-neighboring sites. In a 4-by-4 square lattice, there are a total of 24 and 34 such couplings, respectively.

B. Drive coupling amplitude

To determine the coupling amplitude $g_i = \Omega \times |\alpha_i|$ between qubit i and the common drive, we use a Rabi measurement. We drive each qubit on resonance at $\omega_{\text{com}}/2\pi = 4.5$ GHz, with all other qubits far-detuned to avoid particle exchange. We sweep the drive amplitude A and the drive pulse duration, and measure the Rabi rate at each drive amplitude (Fig. S8a). Next, we find the Rabi rate for different drive amplitudes and use a linear fit to extract g_i/A (Fig. S8b). We repeat this procedure for all qubits in our lattice to find the Rabi rate dependence on the common drive amplitude. By asserting that $|\alpha| = 1$ for the qubit with the strongest coupling to the common drive (qubit 11), we report the relative coupling magnitudes for all the lattice qubits in Fig. S8c (see Tab. S4 C for the values). The common drive coupling strength Ω is therefore related to the drive amplitude A through the linear relationship

$$\Omega(A) = A \times 2\pi \times 49.82(\text{MHz}). \quad (14)$$

We notice that when driving all 16 qubits simultaneously while on resonance, we need to apply a +7.5% correction to the relation above in order for all of our numerical simulations to match the experimental data.

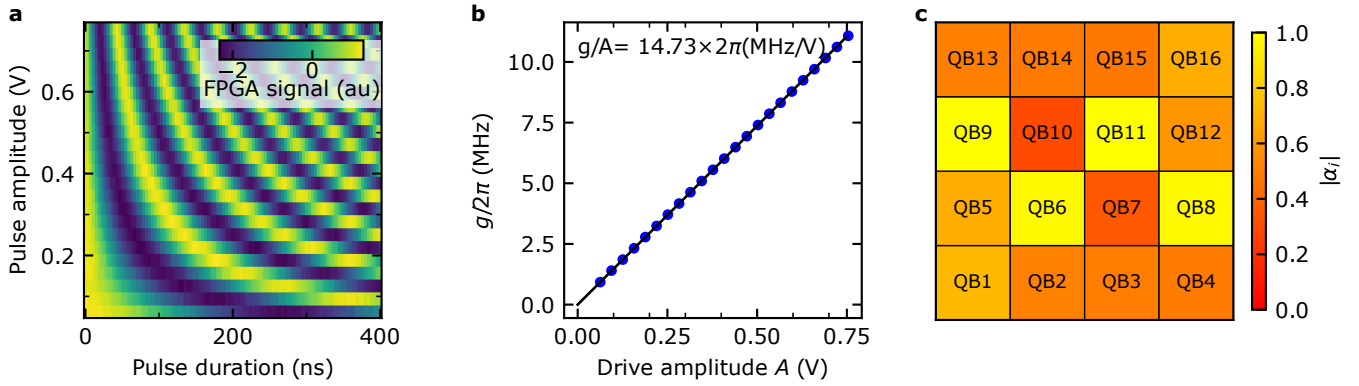


FIG. S8. **Characterizing the drive coupling amplitudes.** (a) Rabi oscillations observed by sweeping the drive pulse duration and amplitude. (b) Rabi rates extracted for different values of the pulse amplitude. (c) The amplitude of the relative drive coupling $|\alpha_i|$ for different qubits in the lattice.

C. Drive coupling phase

We use a two-step approach to characterize the coupling phase $\phi_i = \arg \alpha_i$ between the common drive and each qubit. We need to calculate 15 coupling phases ϕ_i relative to ϕ_1 (the global phase is irrelevant). Single-qubit experiments cannot provide information on coupling phases, so to determine them, we need to conduct drive coupling characterization experiments on multi-qubit subsystems of the entire lattice.

We initially characterize the relative drive coupling phase between different pairs of neighboring qubits using an interferometric approach. For such experiments, we isolate a qubit pair by far-detuning the rest of the lattice. We begin by biasing qubit “A” at $\omega_{\text{com}}/2\pi = 4.5$ GHz and applying a $\pi/2$ -pulse through the common line while detuning qubit “B” to prevent them particle exchange. After this step, the two-qubit system is in the state

$$(|0\rangle - ie^{i\phi_A} |1\rangle) \otimes |0\rangle,$$

where ϕ_A is the phase of the common drive coupled to qubit “A”. Afterward, we bring qubit “B” on resonance with qubit “A”, allowing the two qubits to undergo an iSWAP interaction for time t . After a full iSWAP period, the state of the system is

$$|0\rangle \otimes (|0\rangle - e^{i\phi_A} |1\rangle).$$

Next, we detune qubit “A”, while keeping qubit “B” at ω_{com} and apply a $\pi/2$ -pulse through the common line to qubit “B” resulting in the state

$$|0\rangle \otimes \left(\frac{1 - ie^{i(\phi_A - \phi_B)}}{2} |0\rangle - \frac{ie^{i\phi_B} + e^{i\phi_A}}{2} |1\rangle \right),$$

where ϕ_B is the phase of the common drive coupled to qubit “B”. By sweeping the phase $\Delta\phi$ of the drive and measuring the population on qubit “B” we can find $\phi_A - \phi_B$. The excited state population of the qubit will depend on $\Delta\phi$

$$p(|1\rangle) = \frac{1 + \sin[(\phi_A - \phi_B) - \Delta\phi]}{2}.$$

Hence, when $\Delta\phi = (\phi_A - \phi_B)$ the measured population is $p(|1\rangle) = 1/2$ with a negative slope. We show the pulse sequence for this measurement in Fig. S9a, and representative measurement results using qubits 9 and 10 in Fig. S9b. The gray dashed line indicates $\Delta\phi = (\phi_9 - \phi_{10})$.

Using these relative phases measured between the different qubit pairs, we calculate an absolute phase for each qubit with reference to the drive phase on qubit 1, with $\phi_1 = 0$.

We then take a second step to improve the accuracy of characterizing the drive coupling phases. We isolate different 2×2 plaquettes within our lattice and set the corresponding qubit frequencies to ω_{com} . We apply a weak common drive with strength $\Omega = J/2$ and measure the population on each qubit within the plaquettes as a function of time.

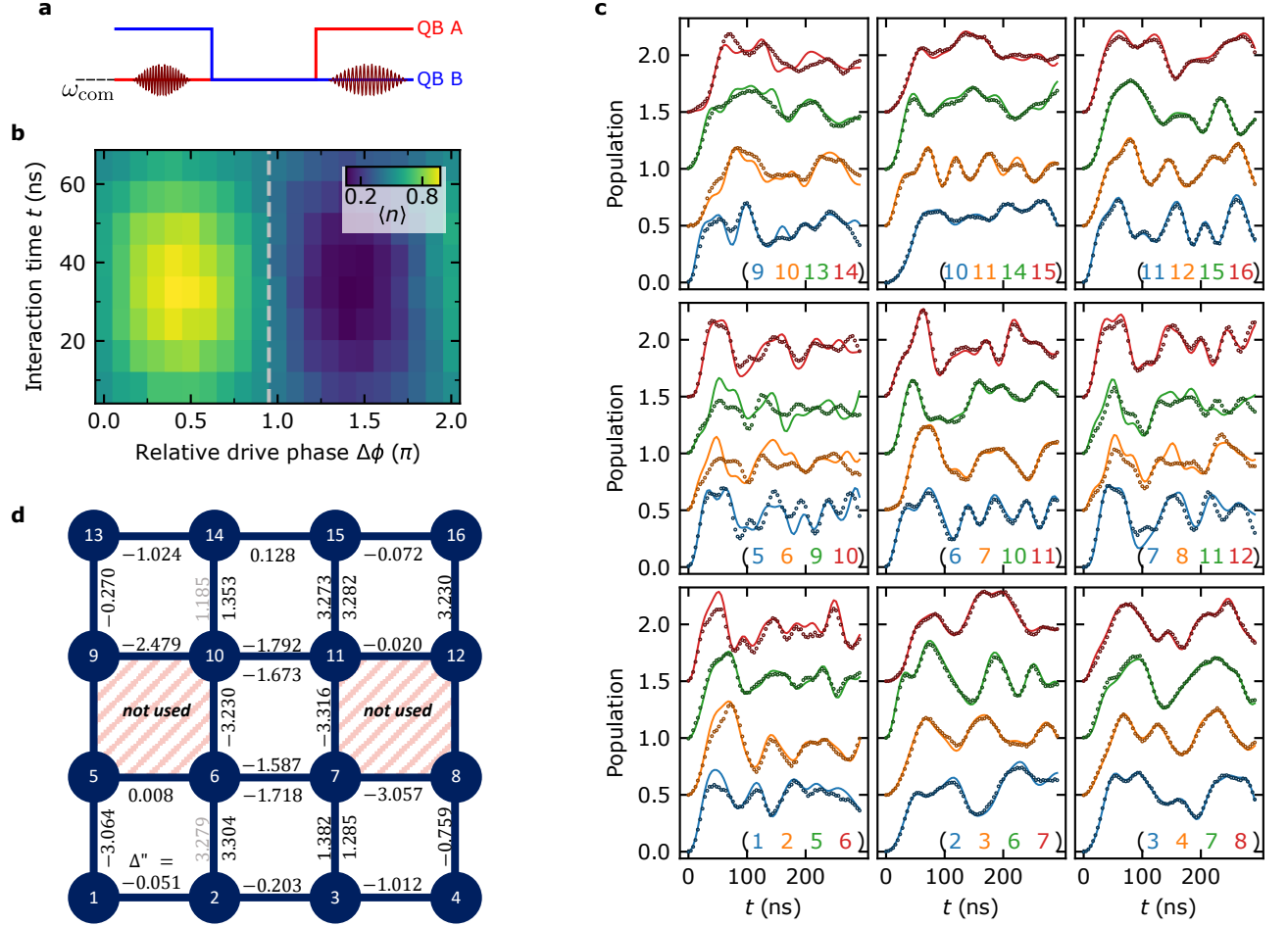


FIG. S9. **Characterizing the drive coupling phase.** (a) Pulse sequence used for interferometric pairwise relative drive coupling phase characterization. (b) Measured population on qubit “B” as a function of sweeping the qubit-qubit exchange interaction time t and the relative drive phase $\Delta\phi$ of the $\frac{\pi}{2}$ -pulse applied to qubit “B” before the readout. The gray dashed line indicates the relative common drive phase between the two qubits. (c) Population versus time for the 9 different 2×2 plaquettes; experimental data are shown as circles and the numerical simulations with the optimized drive coupling phases are shown as lines with colors corresponding to the qubit indices shown in parentheses. (d) The phase difference between the adjacent qubits extracted from the respective 2×2 plaquette fits. The two numbers in grey were unused because those bonds were captured through a higher-quality fit in the neighboring plaquette. The two red-shaded plaquettes were unused because of poor fit quality.

By comparing the measured qubit populations to numerical simulations, we fine-tune the drive coupling phases. To accomplish this, we define the cost function

$$\mathcal{C}(\vec{\phi}) = \sum_{i=1}^4 \sum_{t=0}^{t_f} \left(P_i^{\text{data}}(t) - P_i^{\text{sim}}(t, \vec{\phi}) \right)^2, \quad (15)$$

where $P_i^{\text{data}}(t)$ is the measured population of qubit i after time t and $P_i^{\text{sim}}(t, \vec{\phi})$ is the corresponding simulated value, as parameterized by the three relative drive coupling phases $\vec{\phi}$ in the plaquette. Starting with the characterized drive coupling phases using the pairwise interferometry experiments, we use a gradient-descent-based optimizer to find phase values that minimize the cost function.

We utilize the JAX [16] and OPTAX [17] packages in Python to perform simulations and optimization. We leverage automatic differentiation to compute gradients and determine parameter updates with an Adam optimizer. The initial values provided by the first drive coupling phase calibration step are crucial for gradient-descent-based minimization of the cost function since $\mathcal{C}$ is non-convex in $\vec{\phi}$. We report the experimental data from the different plaquettes and the numerical simulations with the optimized drive coupling phases in Fig. S9c. Generally, we find an excellent agreement

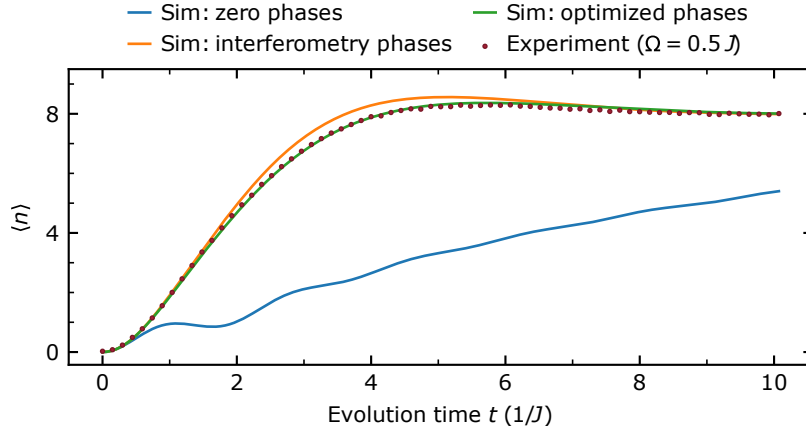


FIG. S10. **Experimental data comparison to simulation.** with (blue line) all phases are set to zero, (orange line) using phases characterized with the pairwise interferometry experiments, and (green line) with fine-tuned phase values.

between the data and numerical simulation. However, for two 2×2 plaquettes, formed by qubits (5, 6, 9, 10) and qubits (7, 8, 11, 12), we observed poor convergence of the optimizer, hence, we exclude the results from those two plaquettes from estimating the phases.

The second step of characterization overdetermines the relative phases: we performed the second step for seven 2×2 plaquettes. Each such measurement yields three independent phase parameters, producing a total of 21 extracted relative phases. Yet there are only 15 physically independent phase values. In order to estimate the most likely drive coupling phases, we take the quality of each fit for each plaquette into consideration (see Fig. S9d). In Tab. S4C we report the final characterized values for $\phi_i = \arg(\alpha_i)$ that we use for the numerical simulations of our 4×4 lattice.

To emphasize the importance of accurately characterizing the drive coupling phase to numerically simulate the behavior of our driven lattice, we consider the time dynamics of the average number of excitations $\langle n \rangle$ in the lattice under a resonant drive. In Fig. S10 we compare experimental data for drive strength $\Omega = J/2$ with numerical simulations with (i) all phases set to zero, (ii) using phases characterized with the pairwise interferometry experiments and (iii) with fine-tuned phase values. We observe that with the fine-tuned phases the numerical simulations of our 4×4 lattice are in excellent agreement with our experimental results.

Qubit index	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16
$ \alpha $	0.73	0.51	0.49	0.49	0.68	0.97	0.35	0.97	1.0	0.3	1.0	0.59	0.5	0.48	0.47	0.68
$\arg(\alpha)$	0.0	6.22	0.1	5.38	3.23	3.26	1.41	4.63	2.88	0.45	3.96	4.0	2.54	1.44	1.14	1.01

TABLE SIII. **Characterized common drive coupling.** amplitudes ($|\alpha|$) and phases ($\arg(\alpha)$).

S5. HARD-CORE BOSE-HUBBARD MODEL ENERGY SPECTRUM

We study the energy spectrum of the hard-core Bose-Hubbard Hamiltonian $\hat{H}_{\text{HCBH}}$ with uniform site energies ($\epsilon_i = 0$) described in Eq. 4 of the main text. We use the measured qubit-qubit coupling strengths, including nearest neighbor and next-nearest neighbor exchange interactions, characterized at 4.5 GHz. In order to efficiently diagonalize the Hamiltonian of our 16-qubit lattice, we first project the $2^{16} \times 2^{16}$ matrix into a subspace spanned by a fixed particle number n . This projection can be accomplished using a $2^{16} \times \binom{16}{n}$ matrix U , where the columns of U are all the permutations of the vectors corresponding to the product states with exactly n particles. The Hamiltonian corresponding to the n particle subspace is therefore calculated through

$$\hat{H}'_n = U^T \hat{H}_{\text{HCBH}} U, \quad (16)$$

where $\hat{H}'_n$ is an $\binom{16}{n} \times \binom{16}{n}$ matrix. By diagonalizing $\hat{H}'_n$ we obtain the eigenenergy distribution for different particle numbers n as shown in Fig. S11a. The largest reduced Hamiltonian is $\hat{H}'_{n=8}$ (describing the subspace of $n = 8$) and its width is $2^{16}/\binom{16}{8} \approx 5$ times less than the full Hamiltonian. Therefore, we expect over $100\times$ speed-up when

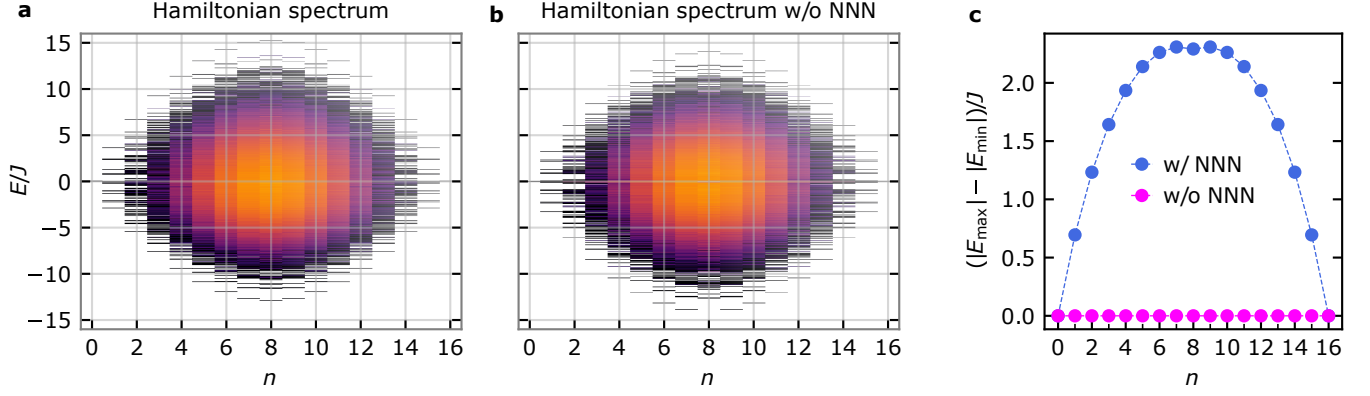


FIG. S11. **Energy spectrum comparison.** The energy spectrum of the characterized system Hamiltonian (a), and the Hamiltonian without the next-nearest coupling (NNN) terms (b). The NNN couplings cause the spectrum to skew towards higher positive energies. (c) The difference between the magnitude of the highest eigenenergy $E_{\max}$ and the lowest eigenenergy $E_{\min}$ for the Hamiltonian with and without NNN terms.

diagonalizing the Hamiltonian of each particle number subspace independently compared to diagonalizing the full system Hamiltonian.

We notice a mild skew towards the higher positive energy eigenenergies in the energy spectrum of our lattice obtained numerically (see Fig. S11a). In contrast, the eigenenergies of the system Hamiltonian excluding the NNN coupling terms are symmetric around zero energy (Fig. S11b). In Fig. S11c we show the difference between the magnitude of the highest eigenenergy $E_{\max}$ and the lowest eigenenergy $E_{\min}$ for the different particle-number subspaces. We notice that for all particle-number subspaces, the eigenenergy distribution is skewed in the positive direction for our system Hamiltonian, which includes NNN couplings, whereas in the absence of NNN coupling, the eigenenergy distribution is symmetric around $E = 0$. Therefore, we conclude that the NNN coupling present in our lattice causes a skew in the eigenenergies of the system.

S6. COHERENT-LIKE STATES ACROSS THE SPECTRUM

For each constant-particle-number subspace of the HCBH Hamiltonian, we observe a variation in the geometric entanglement from the edge to the center of the spectrum [18]. The states at the center of the energy band exhibit a distinct Page curve, while the entropy of the states at the edge of the energy band shows a weak dependence on volume. The geometric entanglement behavior is consistent across different particle-number subspaces, allowing us to probe the entanglement scaling across the many-body spectrum using a superposition of different eigenstates.

Instead of preparing specific eigenstates, we can use a superposition of eigenstates in different regions of the energy spectrum to probe the entanglement properties in a many-body system. We accomplish this by applying a weakly driving a uniform lattice with some detuning δ from the frequency of the qubits.

The hard-core Bose-Hubbard Hamiltonian for a lattice with N sites can be represented in the eigenstate basis as

$$\hat{H}_{\text{HCBH}}/\hbar = \sum_n^N \int d\epsilon \rho(n, \epsilon) (\omega_q + \epsilon) |n, \epsilon\rangle \langle n, \epsilon|, \quad (17)$$

where $|n, \epsilon\rangle$ is the eigenstate and $\rho(n, \epsilon)$ is the density of states with n particles and energy ϵ . In this basis, the driving operator $\hat{\Sigma}$ can be represented as

$$\hat{\Sigma} = \sum_n \int d\epsilon d\epsilon' \left(\rho(n+1, \epsilon) \rho(n, \epsilon') \langle n+1, \epsilon | \hat{\Sigma} | n, \epsilon' \rangle \right) |n+1, \epsilon\rangle \langle n, \epsilon'|. \quad (18)$$

For a weak drive, the driving operator will couple eigenstates that are separated in energy by detuning $\epsilon' - \epsilon = \delta$. Therefore, we can approximate the operators as a combination of energy-raising operators

$$e^{-i\omega_d t} \hat{\Sigma} \approx \int d\epsilon e^{-iH_{\text{HCBH}} t} \hat{A}_\epsilon e^{iH_{\text{HCBH}} t}, \quad (19)$$

with

$$\hat{A}_\epsilon = \sum_n \left(\sqrt{\rho(n+1, \epsilon_{n+1})\rho(n, \epsilon_n)} \langle n+1, \epsilon_{n+1} | \hat{\Sigma} | n, \epsilon_n \rangle \right) |n+1, \epsilon_{n+1}\rangle \langle n, \epsilon_n| \quad (20)$$

where $\epsilon_n = \epsilon \times \delta$. Each $\hat{A}_\epsilon$ couples states on the Hamiltonian spectrum shown in Fig. S11a that are on a line with slope δ and intersection $n = 0$ at energy ϵ . For a lattice initialized with no particles, the only relevant operators are $\hat{A}_0$ and $\hat{A}_0^\dagger$. Therefore, for times t much shorter than $1/\Omega$, the state of the system is approximately

$$|\psi(t)\rangle \approx e^{-i\hat{H}_{\text{HCBH}}t} e^{-i\Omega(\hat{A}_0 + \hat{A}_0^\dagger)t} |\text{vacuum}\rangle \quad (21)$$

which is similar to a multi-mode coherent state. Due to exchange and on-site interactions, the analogy to coherent states becomes weaker at longer times as the number of photons becomes larger.

S7. IMPACT OF THE DRIVE COUPLING PHASE ON EXPERIMENTS

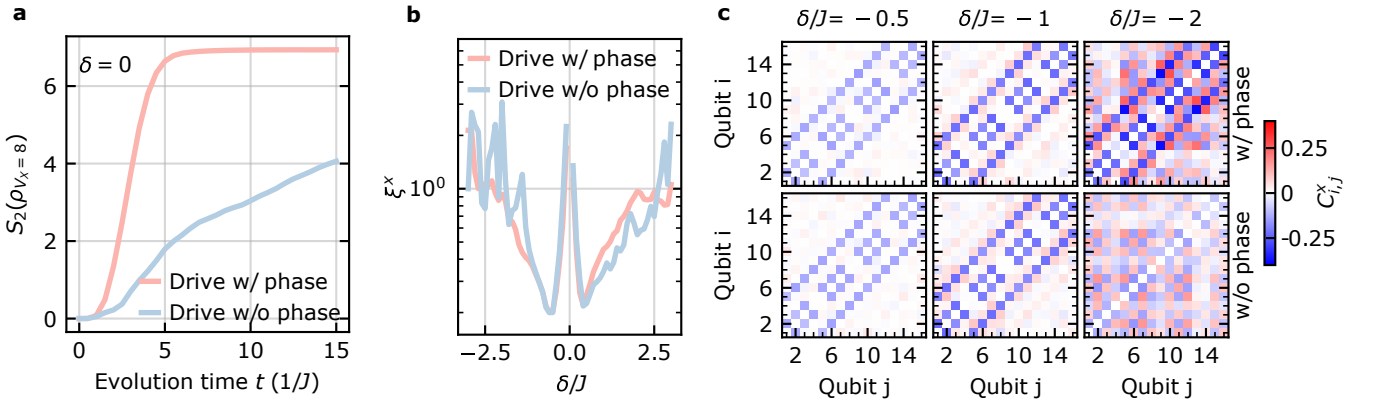


FIG. S12. **Impact of the drive coupling phase on correlations and entanglement.** Simulations of (a) the average entanglement entropy of the 8-qubit subsystems, (b) the correlation lengths, and (c) the correlation matrices within the 4×4 lattice when driven on-resonance with and without the characterized drive phase.

In this section, we discuss the impact of the characterized drive coupling phases on our experiments using numerical simulations. First, we consider the impact of the drive coupling phase on the formation of entanglement entropy within a driven lattice. In Fig. S12a we show the simulated time evolution of the average entanglement entropy of the 8-qubit subsystems within the 4×4 lattice when driven on-resonance with drive strength $\Omega = J/2$. We see that the entropy within the lattice driven with the characterized phases reaches equilibrium much faster than a driven lattice with uniform phases.

Next, we study the correlation lengths within the states prepared with and without the characterized drive phases (Fig. S12b). To ensure that the prepared states are in equilibrium, we simulate evolution of the lattice for $t = 10/J$ under the drive with the characterized coupling phase values, and $t = 50/J$ for the drive with uniform phases. We observe that coherence lengths follow the same trend in both preparation scenarios, with some minor deviations. Finally, we simulate the effect of the drive phase on the correlation matrix in Fig. S12c. The correlation matrices exhibit the same general pattern, however, the correlations between neighboring sites are stronger for states that are prepared using the drive with the characterized coupling phases.

S8. MEASUREMENT OF GLOBAL ENTANGLEMENT

In order to quantify the total entanglement formed, we use the normalized average purity of all the qubits in the lattice $E_{\text{gl}} = 2 - \frac{2}{N} \sum_i \text{Tr}(\rho_i^2)$, which serves as a global entanglement metric [19]. This quantity can be measured with low experimental overhead. We measure E_{gl} by reconstructing the density matrix for each qubit in the lattice via single-qubit tomography. In Fig. S13a we report the measured E_{gl} after driving the lattice for time t with a drive detuning δ . In Fig. S13b, the close agreement between experimental data and numerical simulations, which do not

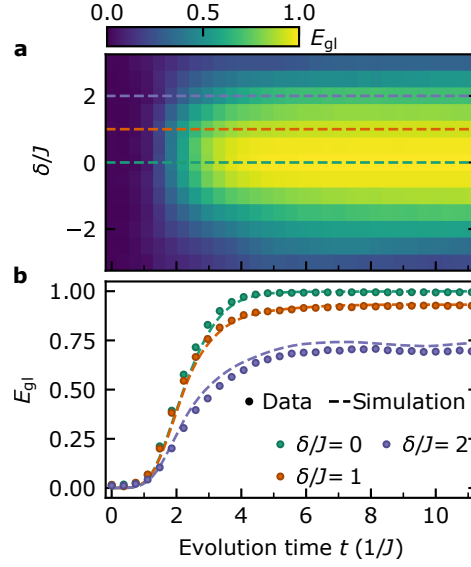


FIG. S13. **Entanglement formation time dynamics.** (a),(b) The global entanglement build-up in the lattice after driving for time t with strength $\Omega = J/2$ and detuning δ from the lattice frequency. Simulations do not include decoherence.

take into account decoherence, suggests that the entanglement formed is between different sites of the lattice, rather than with the uncontrolled environmental degrees of freedom related to decoherence.

While we use E_{gl} to probe the formation of entanglement within our lattice, we cannot distinguish the nature of the entanglement (i.e. whether it is short-ranged or long-ranged). The dynamics of the entanglement formation in the driven lattice reach a steady state after $t \approx 10/J$. We note that as the drive detuning δ gets larger, the steady state value for E_{gl} decreases. This decrease can be attributed to a reduction in the average number of excitations in the lattice with an increase in δ (see Fig. S25).

S9. TOMOGRAPHY SUBSYSTEMS FOR STUDYING THE ENTANGLEMENT BEHAVIOR

In order to tomographically reconstruct the density matrix for a 6-qubit subsystem, we need to measure $3^6 = 729$ Pauli-strings (e.g. $ZXXYZZ$). We use a maximum-likelihood estimation (MLE) routine implemented in Qiskit [20] to reconstruct the density matrix using the measured Pauli-strings. With the aid of our high-fidelity simultaneous single-qubit gates and readout, we can efficiently measure the Pauli-strings needed to reconstruct the density matrix for multiple 6-qubit subsystems using a single round of measurements. Our approach involves initially listing the complete collection of Pauli-strings needed to reconstruct the 6-qubit subsystem (3, 4, 7, 8, 11, 12). We then pair up each of the qubits in the rest of the lattice with a qubit in the subsystem so that their measurement basis is consistent for every measurement. To illustrate how we match the measurement basis of our qubits, we utilize Fig. S14a, in which we use matching colors to indicate the qubits that have identical Pauli operators in each Pauli-string. Qubits 5 and 10 are excluded from tomographic measurements as applying single-qubit gates to these two qubits causes substantial drive crosstalk affecting the other qubits in the system.

Next, we identify all the subsystems that can be reconstructed using the measured Pauli-string. We include all subsystems of size one through six that are nearest-neighbor connected and in which each qubit possesses a distinct color (according to Fig. S14a), yielding 163 unique subsystems. We report a list of all these subsystems, sorted by size V in Table S9. In Fig. S14b we report the distribution of the area (A_X) and the volume (V_X) of the subsystems used for studying the entanglement scaling behavior in the main text. In Fig. S15, we provide a visualization of the entropy extracted for all subsystems as a function of both area and volume.

In our experiments, we measure the expected value of each Pauli string with 2000 shots. We note that the number of samples we can extract for each Pauli string depends on the subsystem size. To see this, notice that measurements of the 6-qubit Pauli strings $XYZXYX$, $XYZXYX$, and $XYZXYX$ all include measurements of the 5-qubit Pauli string $XYZXY$; therefore, while measuring the former we get three times more measurements of the latter. In general, our measurements yield $2000 \times 3^{6-V}$ samples of each volume V subsystem included in the coloring shown in Fig. S14a.

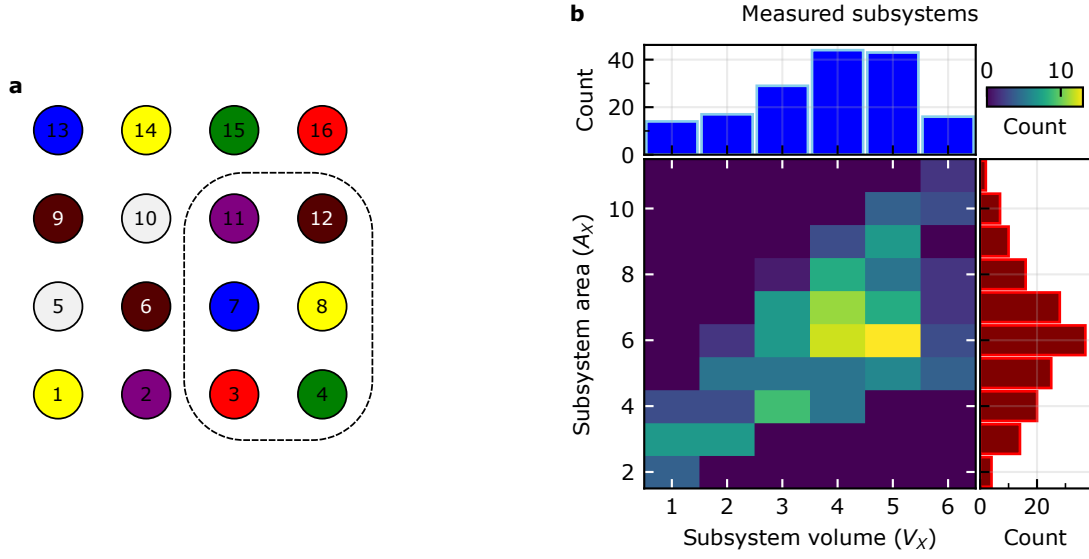


FIG. S14. **Tomography subsystems.** (a) Matching colors indicate the qubits in our lattice that have identical Pauli operators in each measured Pauli-string for tomography. (b) Distribution of the area (A_X) and the volume (V_X) of the measured subsystems.

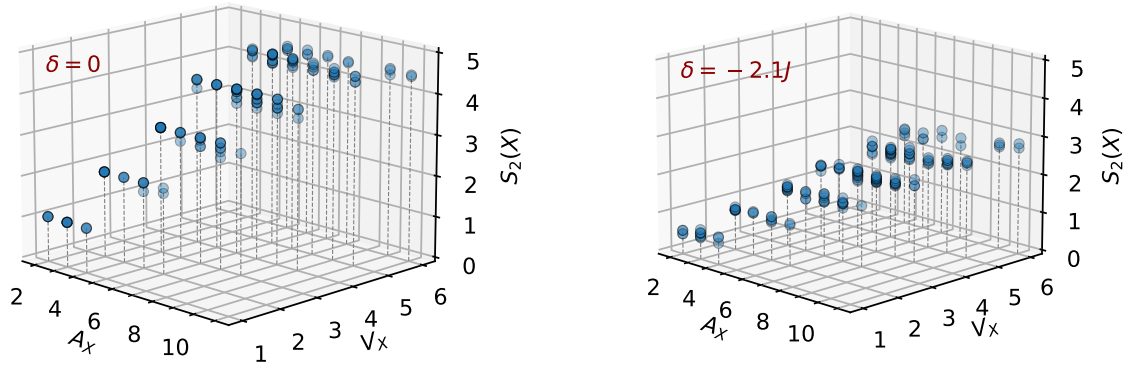


FIG. S15. **Visualizing entanglement in different subsystems.** Scatter plots of the second Rényi entropy extracted of all 163 subsystems, shown as a function of the area (A_X) and the volume (V_X) of each subsystem. Subsystem entropies are shown for the coherent-like states prepared at (left) $\delta = 0$ and (right) $\delta/J = -2.1$.

$V = 1$	(1), (2), (3), (4), (6), (7), (8), (9), (11), (12), (13), (14), (15), (16)
$V = 2$	(1, 2), (2, 3), (2, 6), (3, 4), (3, 7), (4, 8), (6, 7), (7, 8), (7, 11), (8, 12), (9, 13), (11, 12), (11, 15), (12, 16), (13, 14), (14, 15), (15, 16)
$V = 3$	(1, 2, 3), (1, 2, 6), (2, 3, 4), (2, 3, 6), (2, 3, 7), (2, 6, 7), (3, 4, 7), (3, 4, 8), (3, 6, 7), (3, 7, 8), (3, 7, 11), (4, 7, 8), (4, 8, 12), (6, 7, 8), (6, 7, 11), (7, 8, 11), (7, 8, 12), (7, 11, 12), (7, 11, 15), (8, 11, 12), (8, 12, 16), (9, 13, 14), (11, 12, 15), (11, 12, 16), (11, 14, 15), (11, 15, 16), (12, 15, 16), (13, 14, 15), (14, 15, 16)
$V = 4$	(1, 2, 3, 4), (1, 2, 3, 6), (1, 2, 3, 7), (1, 2, 6, 7), (2, 3, 4, 6), (2, 3, 4, 7), (2, 3, 4, 8), (2, 3, 6, 7), (2, 3, 7, 8), (2, 6, 7, 8), (3, 4, 6, 7), (3, 4, 7, 8), (3, 4, 7, 11), (3, 4, 8, 12), (3, 6, 7, 8), (3, 6, 7, 11), (3, 7, 8, 11), (3, 7, 8, 12), (3, 7, 11, 12), (3, 7, 11, 15), (4, 6, 7, 8), (4, 7, 8, 11), (4, 7, 8, 12), (4, 8, 11, 12), (4, 8, 12, 16), (6, 7, 8, 11), (6, 7, 11, 15), (7, 8, 11, 12), (7, 8, 11, 15), (7, 8, 12, 16), (7, 11, 12, 15), (7, 11, 12, 16), (7, 11, 14, 15), (7, 11, 15, 16), (8, 11, 12, 15), (8, 11, 12, 16), (8, 12, 15, 16), (9, 13, 14, 15), (11, 12, 14, 15), (11, 12, 15, 16), (11, 13, 14, 15), (11, 14, 15, 16), (12, 14, 15, 16), (13, 14, 15, 16)
$V = 5$	(1, 2, 3, 4, 6), (1, 2, 3, 4, 7), (1, 2, 3, 6, 7), (2, 3, 4, 6, 7), (2, 3, 4, 6, 8), (2, 3, 4, 7, 8), (2, 3, 4, 8, 12), (2, 3, 6, 7, 8), (2, 3, 7, 8, 12), (2, 4, 6, 7, 8), (3, 4, 6, 7, 8), (3, 4, 6, 7, 11), (3, 4, 7, 8, 11), (3, 4, 7, 8, 12), (3, 4, 7, 11, 12), (3, 4, 8, 11, 12), (3, 6, 7, 8, 11), (3, 6, 7, 11, 15), (3, 7, 8, 11, 12), (3, 7, 8, 11, 15), (3, 7, 11, 12, 15), (3, 7, 11, 14, 15), (4, 6, 7, 8, 11), (4, 7, 8, 11, 12), (4, 7, 8, 12, 16), (4, 8, 11, 12, 16), (6, 7, 8, 11, 15), (6, 7, 11, 14, 15), (6, 7, 11, 15, 16), (7, 8, 11, 12, 15), (7, 8, 11, 12, 16), (7, 8, 11, 15, 16), (7, 8, 12, 15, 16), (7, 11, 12, 14, 15), (7, 11, 12, 15, 16), (7, 11, 14, 15, 16), (8, 11, 12, 15, 16), (9, 11, 13, 14, 15), (9, 13, 14, 15, 16), (11, 12, 13, 14, 15), (11, 12, 14, 15, 16), (11, 13, 14, 15, 16), (12, 13, 14, 15, 16)
$V = 6$	(1, 2, 3, 4, 6, 7), (2, 3, 4, 6, 7, 8), (2, 3, 4, 7, 8, 12), (3, 4, 6, 7, 8, 11), (3, 4, 7, 8, 11, 12), (3, 6, 7, 8, 11, 15), (3, 6, 7, 11, 14, 15), (3, 7, 8, 11, 12, 15), (3, 7, 11, 12, 14, 15), (4, 7, 8, 11, 12, 16), (6, 7, 8, 11, 15, 16), (6, 7, 11, 14, 15, 16), (7, 8, 11, 12, 15, 16), (7, 11, 12, 14, 15, 16), (9, 11, 13, 14, 15, 16), (11, 12, 13, 14, 15, 16)

TABLE SIV. List of reconstructed subsystems.

S10. EXTRACTING s_V AND s_A

After preparing a specific coherent-like state, we extract the volume entropy per site s_V and the area entropy per site s_A using the measured entanglement entropy (S_2) of the different subsystems of our lattice through

$$s_V = \frac{\partial S_2(\rho_X)}{\partial V_X} \Big|_{A_X} \quad (22)$$

$$s_A = \frac{\partial S_2(\rho_X)}{\partial A_X} \Big|_{V_X}. \quad (23)$$

To calculate s_V , we consider the rate of change of $S_2(\rho_X)$ as a function of the subsystem volume V_X for subsystems of constant area (A_X). In Fig. S16a we show the dependence of the average subsystem entropy on V_X for subsystems with a fixed A_X . We fit the slope of the increase in entanglement entropy for each A_X and then calculate s_V as an average of the fitted slopes for each drive detuning value δ (Fig. S16b).

We extract the value of s_A using a similar approach. In Fig. S16c we show the dependence of the average subsystem entropy on A_X for subsystems with a fixed V_X . We fit the slope of the entanglement entropy as a function of the area for each V_X and then calculate s_A as an average of the fitted slopes for each drive detuning value δ (Fig. S16d). Due to the slope being too close to zero, we cannot reliably fit $\partial S_2(\rho_X)/\partial A_X$ in the detuning range $-1 < \delta/J < 1$. In order to avoid divergences in the s_A and s_V fits, we bound the fitted values in the range spanned by $[10^{-4}, 1]$.

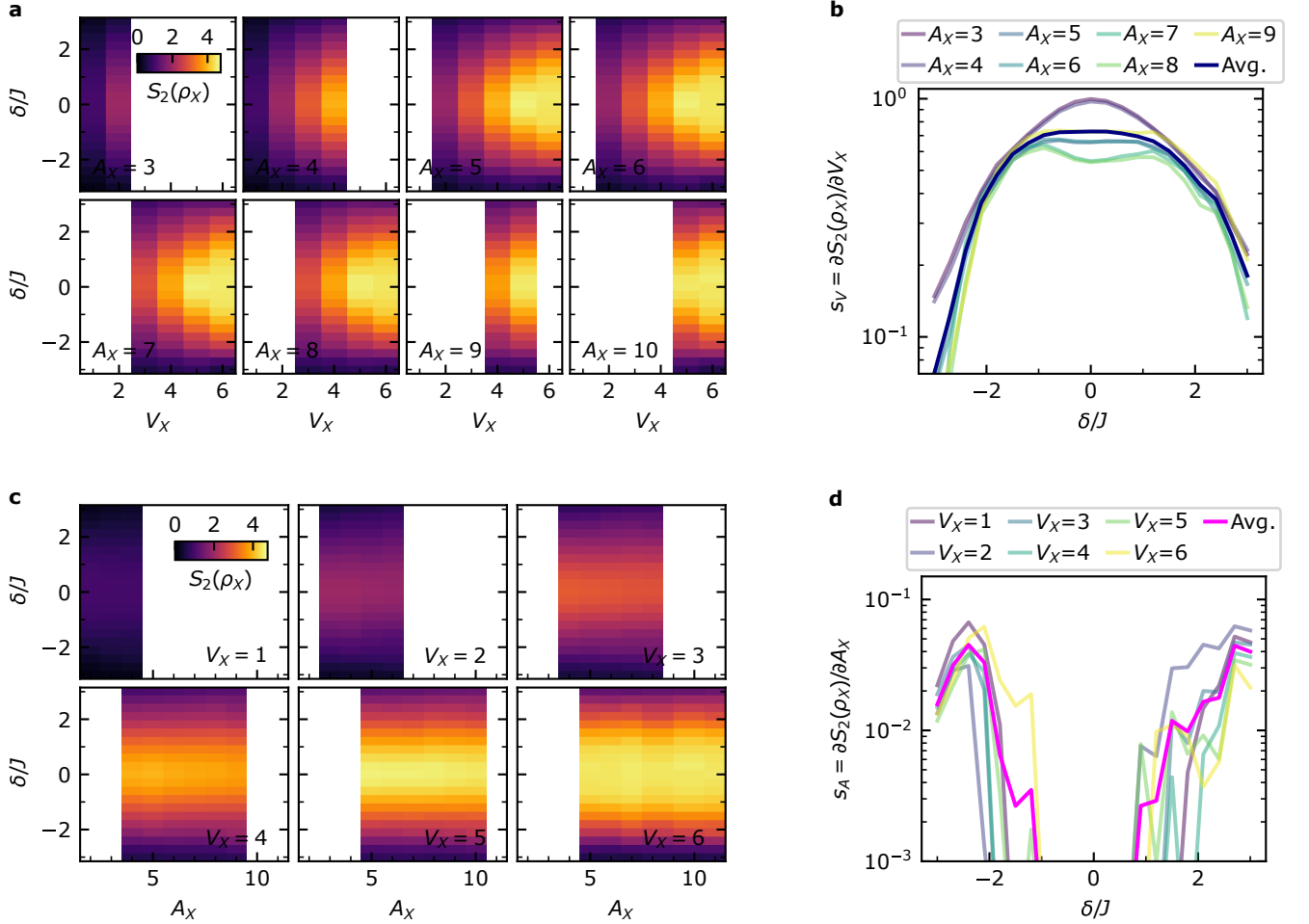


FIG. S16. **Extracting s_V and s_A from tomography data.** (a) Dependence of the average subsystem entropy on V_X for subsystems with a fixed A_X . (b) The volume entropy per site s_V at different drive detunings δ obtained via fitting. (c) Dependence of the average subsystem entropy on A_X for subsystems with a fixed V_X . (d) The area entropy per site s_A at different drive detunings δ obtained via fitting.

S11. EXTRACTING THE ENTANGLEMENT SCALING BEHAVIOR IN LARGER LATTICES

The protocol used in this work to prepare coherent-like states is easily extensible to larger quantum lattices, as is the measurement of correlations. To extract entanglement entropy and Schmidt coefficients, we used full-state tomography of subsystems. Full-state tomography requires exponentially many measurements in system size, so the feasibility of such measurements in larger lattices may not be obvious. In this section, we show that our methods do extend favorably to larger quantum lattices because the size of subsystems that we must consider does not grow with the overall system size.

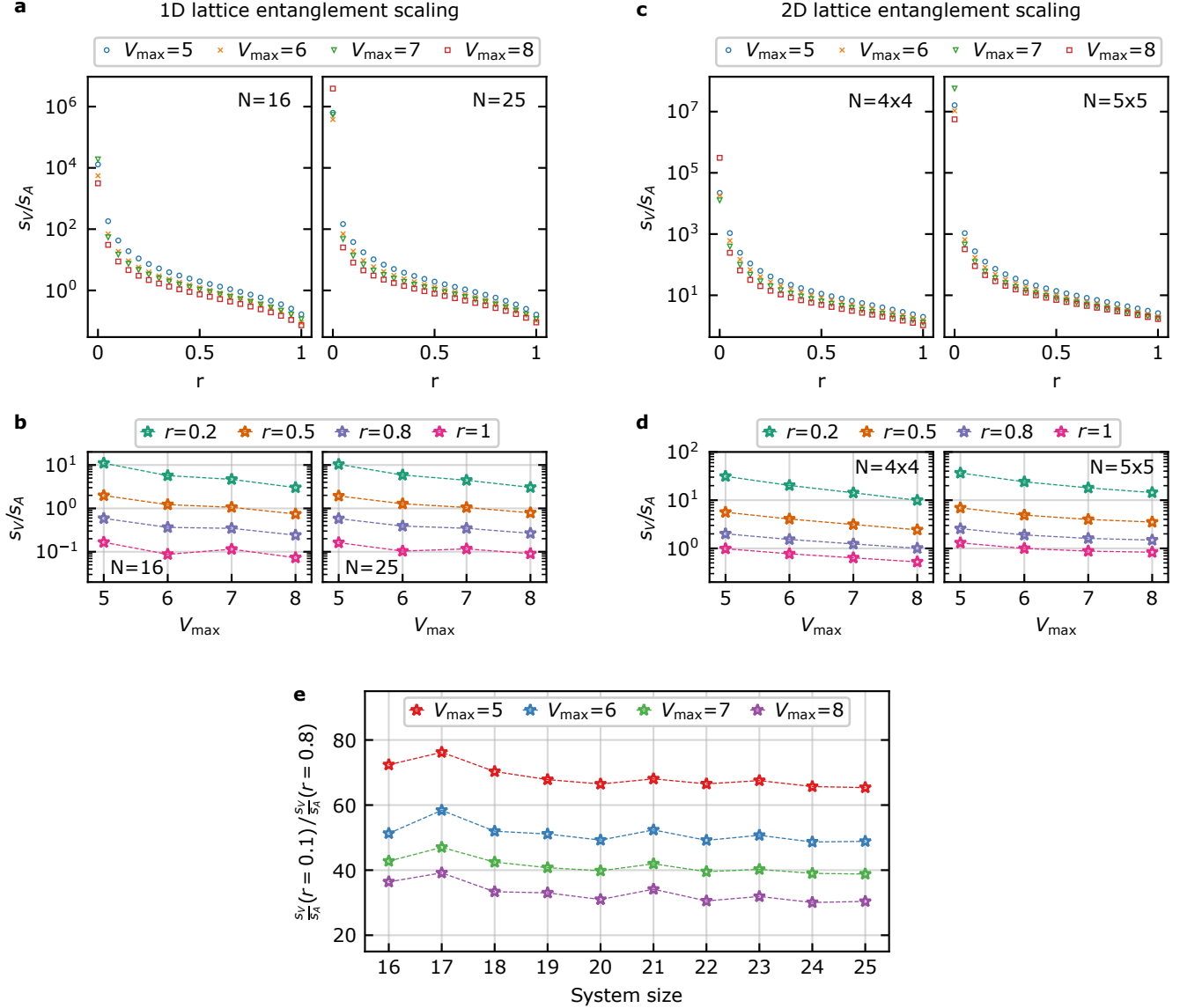


FIG. S17. **Scaling of the experimental approach in extracting the volume-to-area entropy per site ratio s_V/s_A .** (a) The s_V/s_A ratio extracted for various states for 1d HCBH lattices with 16 and 25 qubits. (c) The s_V/s_A ratio variation for different superposition states as a function of the maximum measured subsystem size $V_{\max}$ in 1d lattices. (c) The s_V/s_A ratio extracted for various states for 2d HCBH 4×4 and 5×5 lattices. (d) The s_V/s_A ratio variation for different superposition states as a function of the maximum measured subsystem size $V_{\max}$ in 2d lattices. (e) **The contrast of the s_V/s_A ratio extracted using maximum subsystem size $V_{\max}$ between superposition states with $r = 0.1$ and $r = 0.8$.**

In order to study the geometric scaling of entanglement we do not need to perform full tomography of the entire system. Rather, we can learn the scaling behavior of entanglement within a quantum system by measuring the purity of its subsystems of finite size. In this section, we discuss the scalability of our approach in extracting the volume-to-area entropy per site ratio s_V/s_A to larger lattices.

We begin by considering the quantum state $|\psi\rangle$ in a superposition of a volume-law state $|\psi_{\text{random}}\rangle$ and an area-law state $|\psi_{\text{gs}}\rangle$:

$$|\psi(r)\rangle = \sqrt{1-r}|\psi_{\text{random}}\rangle + \sqrt{r}|\psi_{\text{gs}}\rangle. \quad (24)$$

Here, $|\psi_{\text{random}}\rangle$ is a randomly generated quantum state, which follows volume-law entanglement scaling with high probability [21], and $|\psi_{\text{gs}}\rangle$ is the ground state of the lattice Hamiltonian. By varying the parameter r we can change the participation of the volume-law and area-law states in the superposition, and hence change the geometric entanglement scaling within the state. While fully diagonalizing the Hamiltonian of larger lattices is computationally inefficient, we are able to compute the ground state of HCBH with up to 25 qubits using sparse matrix methods on a commercially available desktop computer.

In Fig. S17a and S17c we numerically study the scalability of our experimental approach in extracting the geometric entanglement scaling in 1D and 2D lattices. We extract the s_V/s_A ratio using the method described in Section S10 for lattices consisting of 16 and 25 qubits by considering subsystems with at most $V_{\text{max}} = 5, 6, 7, 8$ qubits. We observe that the extracted s_V/s_A ratio using different V_{max} values is consistent for the superposition states generated with various values of r . Next, we plot the scaling of the s_V/s_A ratio extracted for superposition states with the four different r factors in Fig. S17b and Fig. S17d for 1d and 2d lattices respectively. We notice there are minor variations in the extracted value as the maximum subsystem size increases, however, the s_V/s_A ratios follow the correct trend (i.e. superposition states with a larger participation of the volume-law state have a higher s_V/s_A).

Finally, to further investigate the capability of our approach in quantitatively distinguishing between area- and volume-law states in large systems, we consider the contrast in the extracted s_V/s_A ratio. In Fig. S17e we show the ratio between the extracted s_V/s_A value for $|\psi(r=0.1)\rangle$ and $|\psi(r=0.8)\rangle$ for 1d lattices with different numbers of qubits. We observe that as the system size becomes larger, the contrast in s_V/s_A values extracted using different maximum subsystem sizes remains consistent. It is worth noting that as V_{max} becomes larger the $\frac{s_V/s_A(r=0.1)}{s_V/s_A(r=0.8)}$ ratio decreases, however, for a given V_{max} the value is consistent for various lattice sizes. This observation suggests that our technique involving measuring the entanglement entropy of subsystems of finite size in a lattice can faithfully distinguish between area- and volume-law states regardless of the overall size of the system.

The numerical simulations discussed in this section confirm that we can study the scaling of entanglement within larger lattices by performing tomography on small subsystems. Furthermore, we note that the size of the subsystems required for studying the entanglement scaling behavior does not depend on the overall lattice size. Using our simultaneous tomographic readout capability we are able to reconstruct the density matrix for various subsystems throughout our lattice by measuring the same number of Pauli strings necessary to reconstruct a single density matrix describing the largest desired subsystem. This feature makes superconducting qubit arrays desirable for studying entanglement within many-body systems.

Lastly, we provide an estimation of the measurement time needed to perform full tomography of a $V = 8$ subsystem. The measurement time is constrained by the time needed to upload measurement sequences from the measurement computer to the measurement hardware, plus the time needed to download measurement outcomes back to the measurement computer. In light of the rapid commercial development of improved measurement hardware, we expect these constraints to be assuaged in the near future. Based on the scaling shown in Fig. ??, a $V = 8$ subsystem requires 3×10^5 measurements for each of the 3^8 Pauli strings, or 2×10^9 total measurements. At the 100 μs measurement period (10^4 samples per second) used in this work, this corresponds to 55 hours of measurement time. Latency associated with measurement sequence upload can be reduced to constant time by programming the field programmable gate arrays onboard measurement hardware to synthesize combinations of pre-compiled pulse waveforms corresponding to each Pauli operator. Assuming state discrimination is performed onboard the measurement hardware, each measurement requires the transfer of 8 bits of data. Data transfer times will then be negligible, assuming 1 Gbit/s communication should become available. Adding extra time for occasional recalibration routines (to correct, for example, for slow drift in $\vec{\Phi}_{\text{offset}}$), we, therefore, expect that near-term hardware developments will enable full tomography of a $V = 8$ subsystem in approximately 3 days.

S12. SCHMIDT COEFFICIENTS SCALING

The Schmidt decomposition is a useful tool for representing entangled states. For a system AB in a pure state $|\psi\rangle$, the theorem states that there exist orthonormal states $|k_A\rangle$ for subsystem A , and $|k_B\rangle$ for subsystem B , referred to as the Schmidt bases, such that

$$|\psi\rangle = \sum_k \lambda_k |k_A\rangle |k_B\rangle, \quad (25)$$

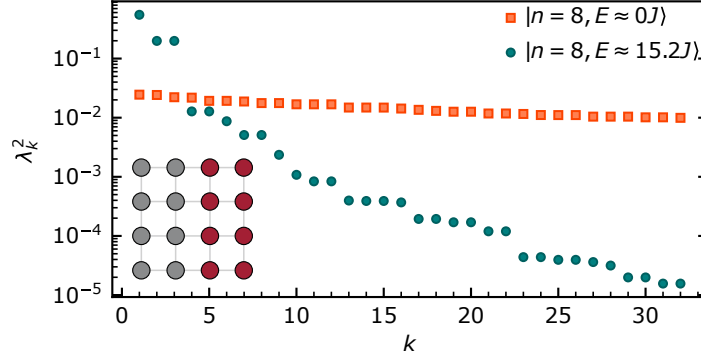


FIG. S18. **Schmidt decomposition (a)** The 32 largest Schmidt coefficients for a bipartition of a volume-law eigenstate at the center (energy $E \approx 0$) of the 16-qubit HCBH energy spectrum with an area-law eigenstate at the edge (energy $E \approx 15.2J$) of the spectrum.

where λ_k coefficients are non-negative real numbers satisfying $\sum_k \lambda_k^2 = 1$, and are known as Schmidt coefficients [22]. The Schmidt coefficients provide insight into the “amount” of entanglement between the two subsystems. For two maximally entangled subsystems, each with n qubits, $\lambda_k = 1/\sqrt{2^n}$, which results in an entanglement entropy $S_2(\rho_A) = n$. The pairwise entanglement for any two qubits in such a system is, therefore, exponentially small in the system size. In slightly entangled states, only a relatively small number of the Schmidt coefficients have a significant contribution to the weight of the state. Therefore, area-law states can be compressed and represented by a truncated Schmidt decomposition [21]

$$|\psi_{\text{trunc}}\rangle = \sum_k^\chi \lambda_k |k_A\rangle |k_B\rangle \quad (26)$$

using the χ largest Schmidt coefficients. Independent of the system size, we can truncate the Schmidt decomposition of an area-law state with a fine number χ and up to an arbitrary error $\epsilon > 0$ such that

$$|| |\psi\rangle - |\psi_{\text{trunc}}\rangle ||^2 < \epsilon. \quad (27)$$

Using the Schmidt decomposition, we can represent the reduced density matrix ρ_A in the Schmidt basis

$$\rho_A = \text{Tr}_B(\rho_{AB}) = \sum_i \lambda_i^2 |i_A\rangle \langle i_A|. \quad (28)$$

Hence, by diagonalizing the measured density matrices of our lattice subsystems we can find the Schmidt coefficients of the decomposition between the different subsystems and their complement. In Fig. S18, we compare the largest 32 Schmidt coefficients for an equal bipartition of a volume-law eigenstate at the center of the 16-qubit HCBH energy spectrum with an area-law eigenstate at the edge of the spectrum. We observe that a small number of Schmidt states contain nearly all the weight of the decomposition for the eigenstate at the edges of the spectrum, whereas, for the eigenstate at the center of the energy spectrum, the Schmidt coefficients are roughly equal to one another. This observation suggests that the area-law states at the edges of the energy band can be efficiently numerically simulated using tensor network methods [23], whereas volume-law states are harder to approximate classically.

Like the entropy, the Schmidt coefficients extracted from reconstructed density matrices are sensitive to the number of measurement samples recorded of each Pauli string in tomography. In Fig. S19, we show the number of Schmidt coefficients needed to decompose, with 0.999 accuracy, the density matrices reconstructed from simulations of sampled measurements discussed in the main text (Methods – Measurement Sampling Statistics section), as a function of number of samples.

S13. ESTIMATING THE IMPACT OF DECOHERENCE ON MEASURED SYSTEM ENTROPY

In the duration of the experiment, our system is subject to decoherence in the form of qubit energy relaxation and dephasing. In our system the average relaxation time is $T_1 = 22.8\mu\text{s}$ and the average dephasing time is $T_\phi = 2.4\mu\text{s}$

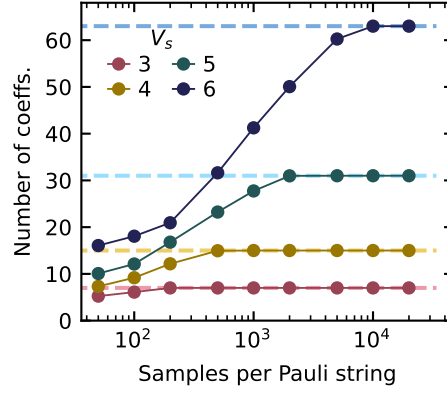


FIG. S19. **Simulation of Schmidt decomposition versus measurement samples.** The dark markers present the number of Schmidt coefficients needed to represent with 0.999 accuracy the density matrices extracted from simulated tomography of subsystems as a function of the number of measurement samples of each Pauli string used for density matrix reconstruction. The light dashed lines represent the number of Schmidt coefficients needed to decompose the exact simulated state with the same fidelity. Maximum-likelihood estimation is used for density matrix reconstruction. The coherent-like state generated at $\delta = 0$ is used for this figure. For each subsystem volume V_s , results were averaged over the same subsystems used in the experiment.

at ω_{com} . To prepare the highly-entangled, coherent-like states, we drive the lattice for $t = 270$ ns, followed by the tomographic readout of each qubit. As all qubits undergo Rabi-like oscillations under the drive, which is near the qubit frequency, the decoherence process closely resembles relaxation during rotary echo [24, 25]. In this case, the dominant decay process is exponential at the rate

$$\Gamma = \frac{3}{4}\Gamma_1 + \frac{1}{2}\Gamma_\nu, \quad (29)$$

where $\Gamma_1 = 1/T_1$ and $\Gamma_\nu = \frac{1}{2}S_z(\Omega)$. Here, the spectral noise power S_z is evaluated at the drive amplitude chosen in experiment $\Omega = J/2$. Assuming $1/f$ -type dephasing noise, $S_z(\Omega) = \frac{A_\Phi}{2\pi\Omega} \left| \frac{\partial\omega}{\partial\Phi} \right|^2$, where A_Φ is the flux noise amplitude. Rotary echo and spin locking measurements were not explicitly performed; we instead extract A_Φ from a measurement of the spin-echo dephasing rate $\Gamma_{\text{echo}} = \sqrt{A_\Phi \ln 2} \left| \frac{\partial\omega}{\partial\Phi} \right|$. We use measurements of QB5 at the common lattice frequency ω_{com} , where $T_1 = 17.6 \mu\text{s}$ and $1/\Gamma_{\text{echo}} = 19.7 \mu\text{s}$, from which we extract an inconsequential value $1/\Gamma_\nu > 1$ ms. The dominant loss channel is then longitudinal decay in the lab frame at timescale $1/\Gamma \approx \frac{4}{3}T_1 = 23.5 \mu\text{s}$.

Consider a subsystem X and its complement $\bar{X}$ within the system S . Defining the error $\gamma = 1 - e^{-\Gamma\tau}$, where $\tau = 280$ ns is the experiment duration, and approximating the system dephasing as a quantum channel to the first order in γ , the system density matrix after undergoing dephasing is

$$\rho' = (1 - n\gamma)\rho + \gamma \sum_i^n \sigma_i^z \rho \sigma_i^z, \quad (30)$$

where $\rho = |\psi\rangle\langle\psi|$ is the density matrix of the pure state of the system with $n = 16$ qubits. We define the reduced density matrix without dephasing $\rho_X = \text{Tr}_{\bar{X}}\rho$. Because $\text{Tr}_X \sigma_z^i \rho \sigma_z^i = \rho_X$ for $i \in \bar{X}$, the reduced density matrix of the subsystem is:

$$\rho'_X = \text{Tr}_{\bar{X}}\rho' = (1 - n_X\gamma)\rho_X + \gamma \sum_{i \in X} \sigma_i^z \rho_X \sigma_i^z. \quad (31)$$

The subsystem purity is then:

$$\text{Tr}(\rho'^2_X) = (1 - n_X\gamma)^2 \text{Tr}[\rho_X^2] + \gamma^2 \sum_{i,j \in X} \text{Tr}[\sigma_i^z \rho_X \sigma_i^z \sigma_j^z \rho_X \sigma_j^z] + 2\gamma(1 - n_X\gamma) \sum_{i \in X} \text{Tr}[\rho_X \sigma_i^z \rho_X \sigma_i^z], \quad (32)$$

or:

$$\text{Tr}(\rho'^2_X) = (1 - n_X\gamma(2 - n_X\gamma - \gamma)) \text{Tr}[\rho_X^2] + \gamma^2 \sum_{i \neq j \in X} \text{Tr}[\sigma_i^z \rho_X \sigma_i^z \sigma_j^z \rho_X \sigma_j^z] + 2\gamma(1 - n_X\gamma) \sum_{i \in X} \text{Tr}[(\rho_X \sigma_i^z)^2]. \quad (33)$$

Notice that the latter two terms are density-density correlators and the squared density expectation values, respectively. In the volume-law limit, both terms vanish [26] (intuitively, this is because volume-law states are similar to infinite temperature states). Approximating these terms as zero, the subsystem entropy with dephasing is then:

$$S_2(\rho'_X) \approx S_2(\rho_X) - \log(1 + n_X^2 \gamma^2 - 2n_X \gamma + n_X \gamma^2). \quad (34)$$

Therefore, the entropy contribution due to dephasing is approximately

$$S_{2,\text{dephasing}} = -\log(1 + n_X^2 \gamma^2 - 2n_X \gamma + n_X \gamma^2). \quad (35)$$

This estimate for the additional entropy due to decoherence is used in Fig. 3c of the main text, and is reproduced here for all subsystem sizes in Fig. S20.

Away from the volume-law limit $\delta/J \lesssim 1$, we can estimate the leading correction to the above approximation $2\gamma(1 - n_X \gamma) \sum_{i \in X} \text{Tr}[(\rho_X \sigma_i^z)^2] \approx 0$, by the mean-field expression $2\gamma n_X \left(\frac{N/2 - \langle n \rangle}{N/2} \right)^2$ where $N = 16$ is the system size, and measurements of $\langle n \rangle$ at various δ are shown in Fig. S25.

While the intended initial state in this experiment is a vacuum state $\langle n \rangle = 0$, in reality there may be finite initial population due to non-zero qubit temperature. As indicated in Table S2, the average thermal population is roughly $\varepsilon_0 = 0.04$. Subsystem initial states therefore have entropy $S_2^{\text{thermal}} \approx -n_X \log(1 - 2\varepsilon)$. However, because the qubits are driven on or near resonance during the experiment, we do not expect S_2^{thermal} to contribute directly to additional entropy of the final state. Instead, an initial thermal excitation can be thought of as initializing the experiment in a superposition of the single-particle states. The experiment will then populate a swath of the many-body spectrum slightly more broad than the ideal coherent-like state, reducing the observed contrast between area-law-like and volume-law states.

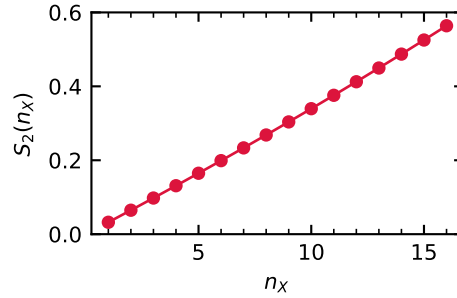


FIG. S20. **Entropy contribution from decoherence.** An estimate of the entanglement entropy contribution to the measurement from spin-locking dephasing (with timescale $1/T$).

S14. TIME DYNAMICS SIMULATIONS

In this section, we report time dynamics simulation results to get a better understanding of the steady-state behavior of a driven lattice. First, we compare the excitation population for a 3×3 and a 4×4 lattice in Fig. S21a. We observe that both systems reach a steady-state population of half-filling by $t = 10/J$ for different drive strengths, suggesting that the steady-state time scale does not strongly depend on the lattice size. However, the 3×3 lattice exhibits small fluctuations in population due to the finite lattice size and deviation from the thermal limit.

Next, we consider the evolution of the subsystem entropy in a driven 4×4 lattice with a drive strength $\Omega = J/2$, and detuning $\delta = 0$ (Fig. S21b). We observe the formation of entanglement within the system as indicated by the increase in entropy. The entanglement entropy of the subsystems of various sizes approaches the value predicted by the Page curve a volume-law state by $t = 10/J$.

In Fig. S21c and S21d we report the simulated time dynamics of s_V/s_A ratios and correlation lengths ξ^x for various drive detunings δ . We notice minor variations in these metrics for drive times exceeding $t = 10/J$. The time evolution simulations in this section suggest that the states prepared with a drive time of $t = 10/J$ in our experiments are a very close approximation to the desired steady-state coherent-like states.

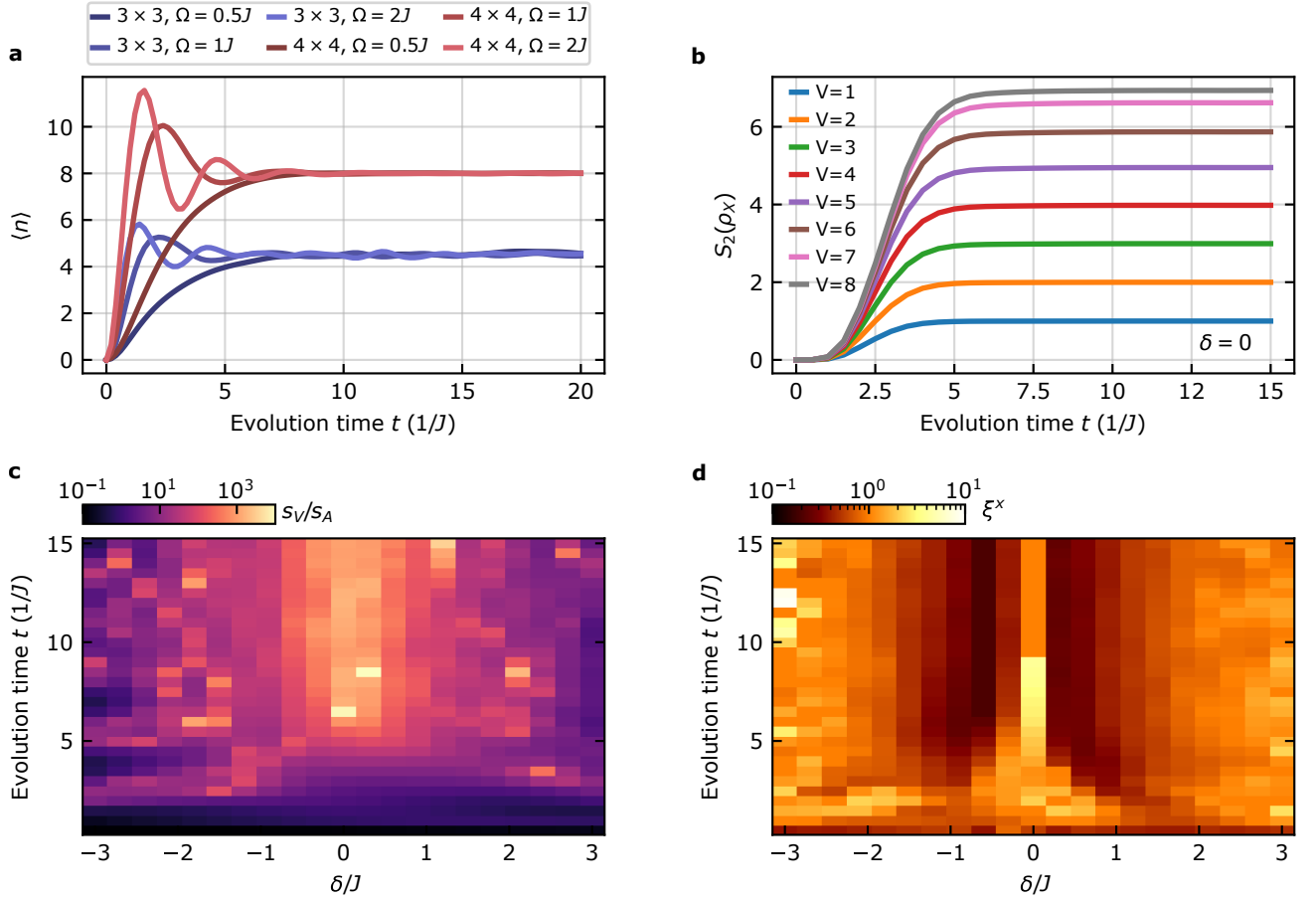


FIG. S21. **Time dynamics simulations.** (a) Excitation population for driven 3×3 and 4×4 lattices. (b) The time evolution of the entropy of subsystems of different sizes in lattice driven with drive strength $\Omega = J/2$ and detuning $\delta = 0$. (c) The time dynamics of s_V/s_A ratios for various drive detunings δ . (d) The time dynamics of correlation lengths ξ^x for various drive detunings δ .

S15. COHERENT-LIKE STATES IN 1D

In two dimensions, the hard-core Bose-Hubbard model is non-integrable, while in one dimension the model is integrable. While a general statement is beyond the scope of the present work, eigenstate thermalization is known to fail in many integrable systems. It is therefore interesting to compare the behavior of the two-dimensional (2D) system presented in this work with that of an analogous one-dimensional (1D) system. In this section, we present numerical simulations of coherent-like states prepared in a 14-qubit one-dimensional hard-core Bose-Hubbard chain, both with and without next-nearest neighbor exchange interactions. These simulations do not include decoherence.

Unlike in two dimensions, in the one-dimensional system, the population does not reach a steady state in time of order $1/J$. The average number of excitation continues to oscillate at for tens of characteristic periods $1/J$, with the specifics of the oscillation amplitude and ringdown depending on δ (see Fig. S22a and S22b). The analogous dynamics in the presence of exaggerated next-nearest neighbor coupling $J_{\text{NNN}} = J/6$ are shown in Fig. S23a and S23b.

In Fig. S22d we present the two-point correlations averaged across all distance-3 qubit pairs in the 1D lattice as a function of time. Unlike in the 2D case, here the correlations fluctuate significantly in time, even after time $10/J$. We cannot extract meaningful correlation lengths because the correlations do not reach a steady state. Correlations at different distances, presented as time-averaged values between $t/J = 15$ and $t/J = 20$, are shown in Fig. S22e.

As in the 2D case, we find the largest entanglement entropy for $\delta = 0$. In 2D, the entropy saturates after a time of order $1/J$ (Fig. S26). In 1D, the entropy continues to grow after many multiples of the characteristic time $1/J$. As shown in Fig. S22c, the continued entropy growth at $t \gg 1/J$ is particularly pronounced when $\delta/J \lesssim 1$ and persists in the presence of weak next-nearest neighbor exchange interactions (Fig. S23c).

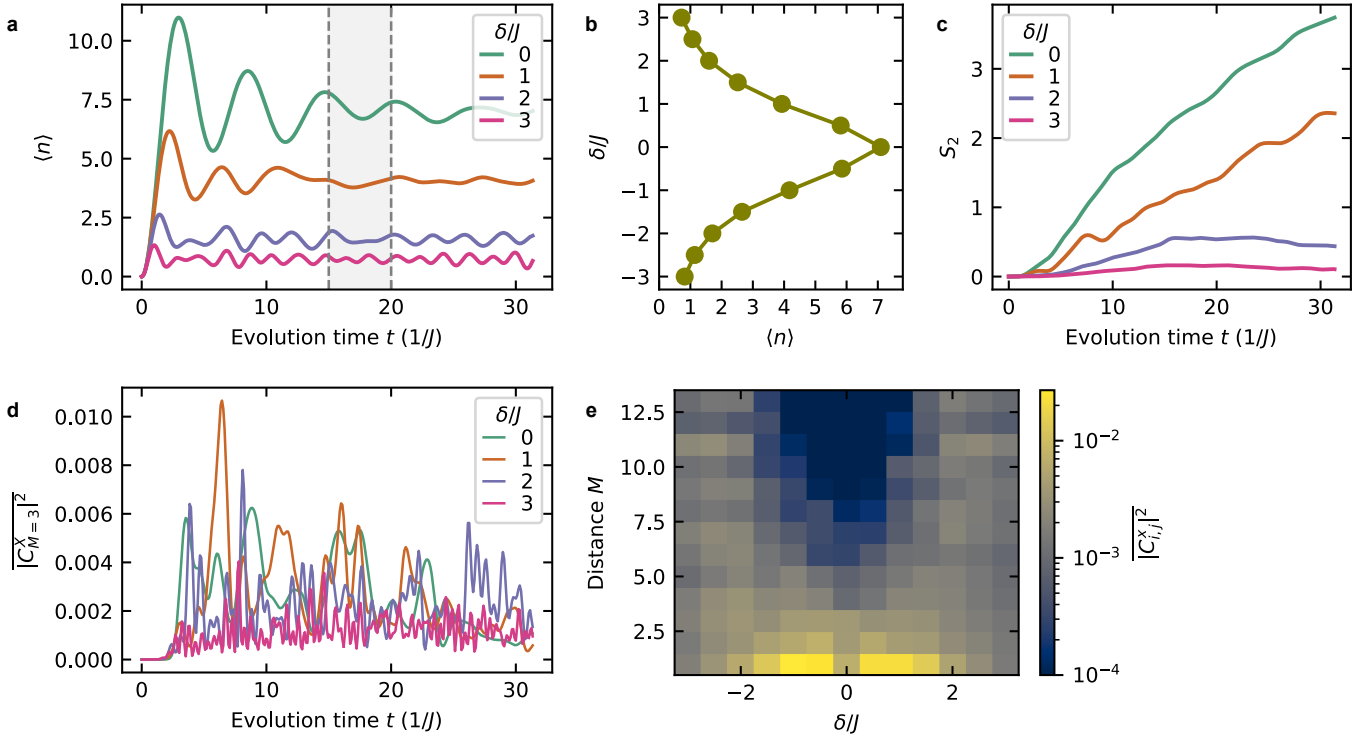


FIG. S22. **Simulations of coherent-like states in a one-dimensional 14 qubit system.** (a) Measured total number of excitations $\langle n \rangle$ while driving the system with drive strength $\Omega = J/2$ and various detunings δ for time t . (b) Total number of excitations versus δ , averaged between $t/J = 15$ and $t/J = 20$. (c) The second Rényi entropy of the subsystem consisting of the first seven qubits in the chain, shown as a function of time for various drive detunings. (d) The 2-point distance-3 correlator $|C_{M=3}^X|^2$, averaged over all distance-3 pairs in the lattice, as a function of time for various drive detunings. (e) The time-averaged 2-point correlator shown as a function of δ and pair distance M . The shaded region in a indicates the arbitrarily-chosen time window throughout which the results shown in subfigures b and e were averaged.

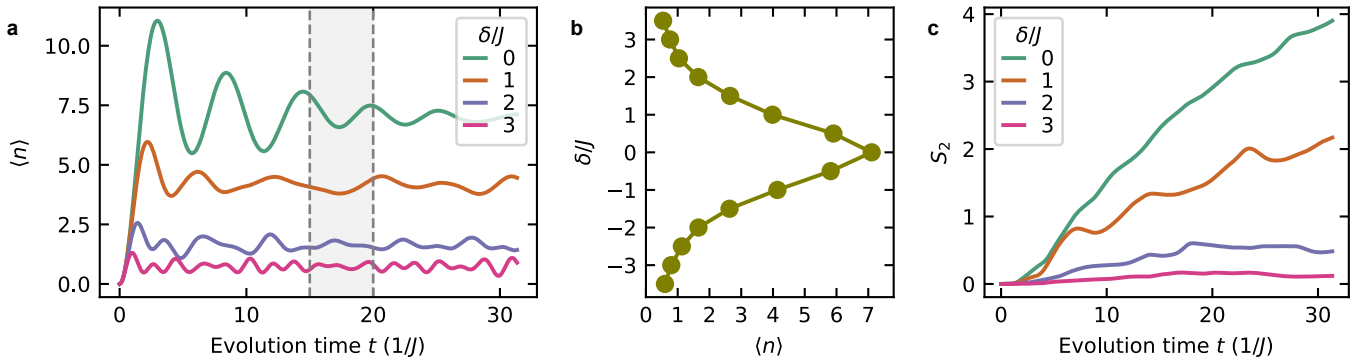


FIG. S23. **Simulations of coherent-like states in a one-dimensional 14 qubit system with a next-nearest neighbor coupling of strength $J_{\text{NNN}} = J/6$.** (a) Measured total number of excitations $\langle n \rangle$ at drive strength $\Omega = J/2$. (b) Total number of excitations versus δ , averaged between $t/J = 15$ and $t/J = 20$ (indicated by the shaded region in a). (c) The second Rényi entropy of the subsystem versus time, for the subsystem consisting of the first seven qubits in the chain.

S16. EXTENDED DATA

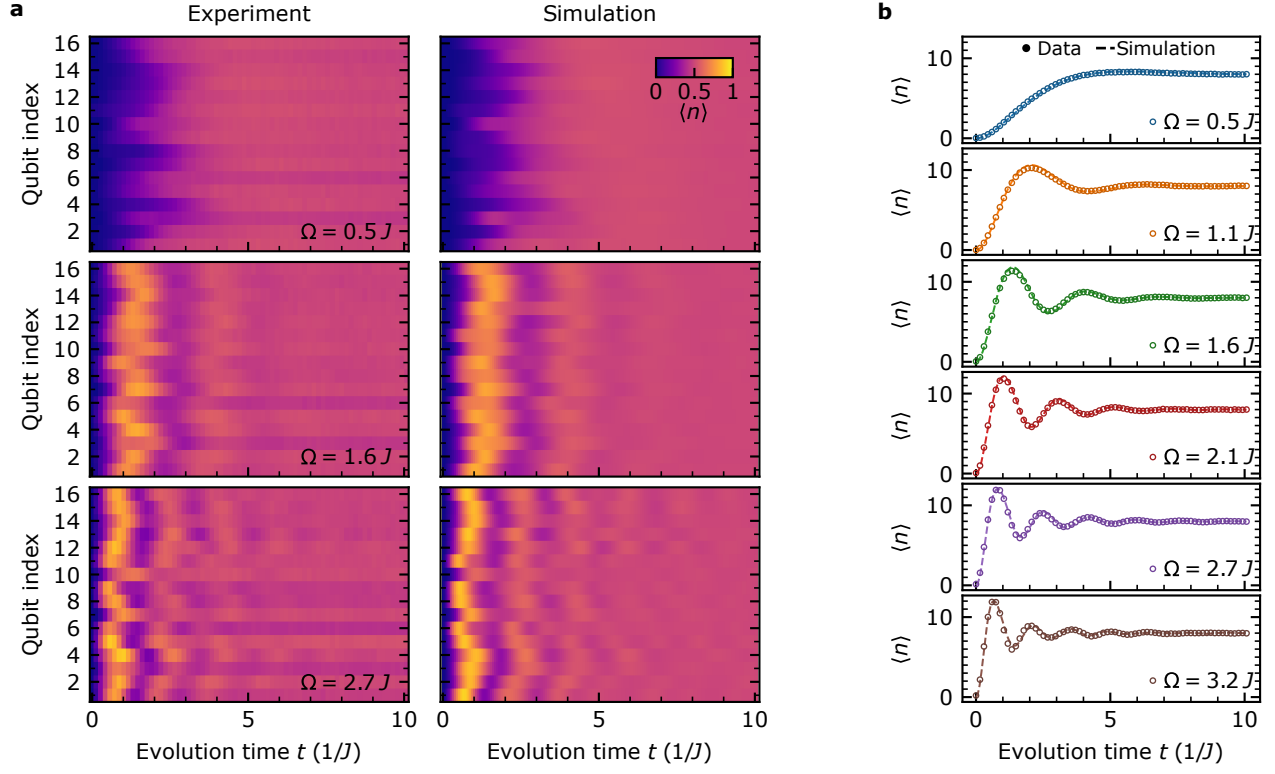


FIG. S24. **Extended data on generating coherent-like states.** (a) Measured population on each qubit in the uniform lattice while driving the system on resonance for time t with different values of drive strength Ω . (b) The total number of excitation $\langle n \rangle$ in the uniform lattice while driving the system on resonance for time t . We observe excellent agreement between experimental data and numerical simulations using the characterized Hamiltonian parameters.

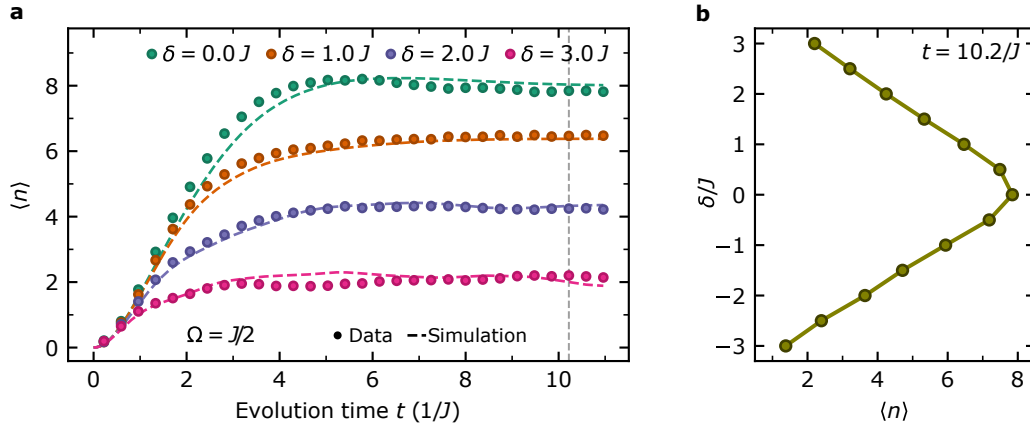


FIG. S25. **Extended data on coherent-like states population with detuned drive.** (a) Measured total number of excitation $\langle n \rangle$ in the uniform lattice while driving the system with drive strength $\Omega = J/2$ and detuning δ for time t . (b) The lattice particle number $\langle n \rangle$ measured in experiments for states at approximately $t \approx 10/J$. The number of excitations at equilibrium will vary depending on δ . When $\delta = 0$, the lattice reaches half-filling, and as the magnitude of the detuning increases, the value of $\langle n \rangle$ at equilibrium decreases.

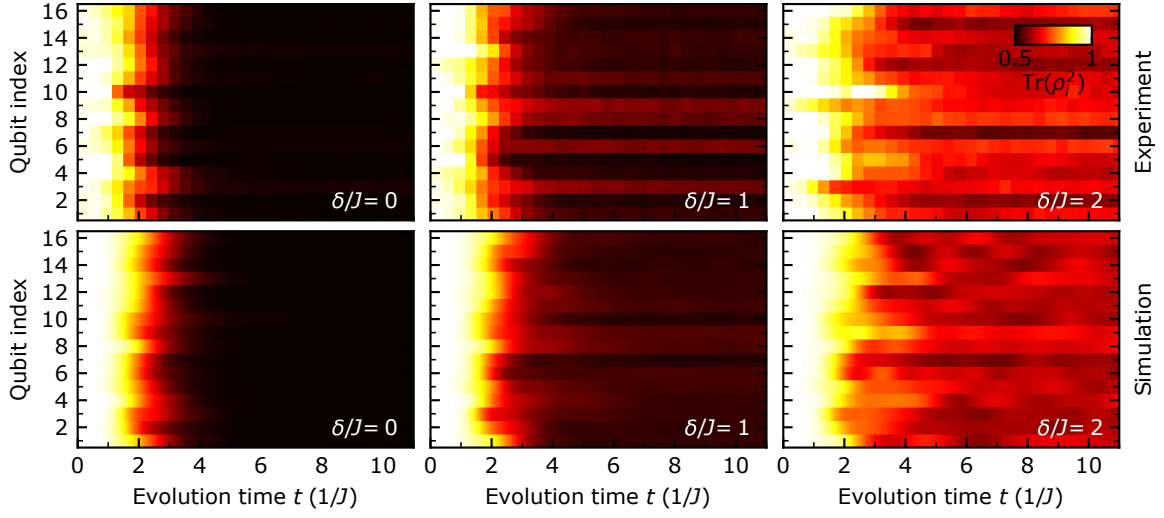


FIG. S26. Extended data on single-qubit purity evolution.

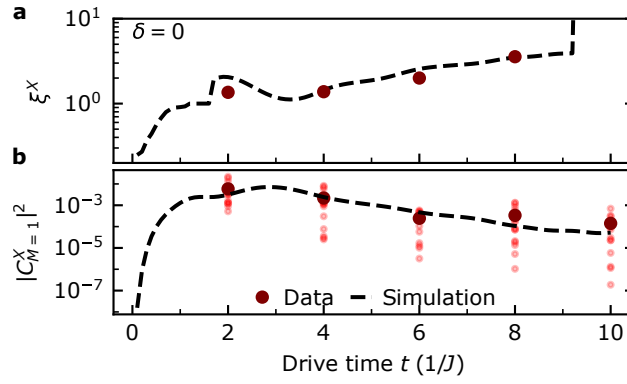


FIG. S27. **Evolution of 2-point correlations** at $\delta = 0$. The evolution of the correlation lengths **(a)** and the 2-point correlator values for neighboring qubits $|C_{M=1}^X|^2$ **(b)** as a function of the duration of the common drive at $\delta = 0$. After $t = 1/J$, the correlation lengths become longer, and the $|C_{M=1}^X|^2$ decreases.

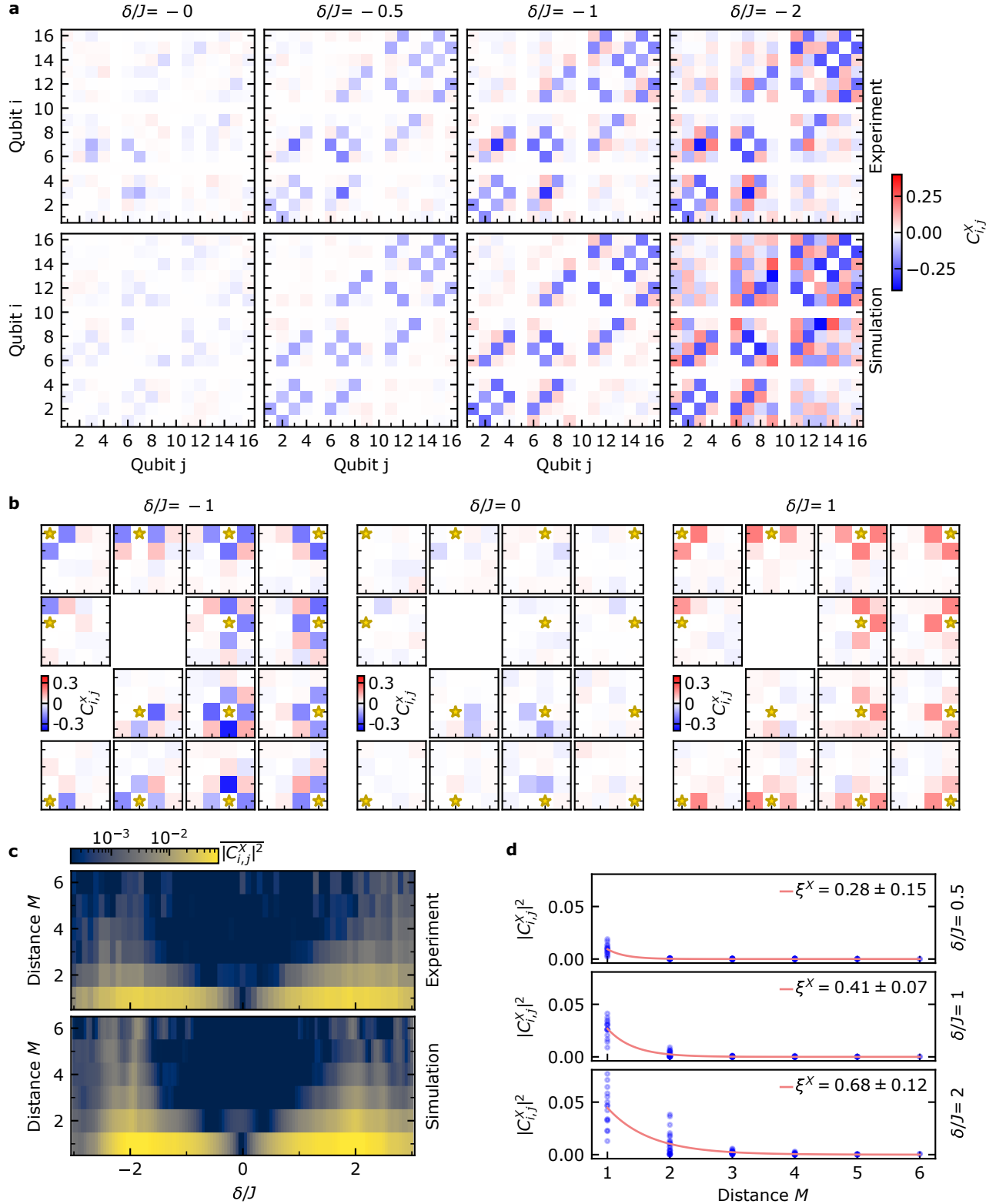


FIG. S28. **Extended data on 2-point correlations.** (a) Correlation matrix constructed using the pairwise correlators $C_{i,j}^X \equiv \langle \hat{\sigma}_i^x \hat{\sigma}_j^x \rangle - \langle \hat{\sigma}_i^x \rangle \langle \hat{\sigma}_j^x \rangle$ between the different qubit pairs in our system for different values of the drive detuning δ . (b) Spatial structure of the measured pairwise correlations at $\delta = -J, 0, J$. We notice that the sign of the drive detuning determines the sign of the correlations between the nearest neighbors. For $\delta < 0$ (> 0), corresponding to population in lower (higher) energy half of the many-body band, we observe antiferromagnetic (ferromagnetic) correlations between nearest neighbors. This observations reflects that $J > 0$ in our device. Correlations among next-nearest neighbors are weaker, reflecting that the coherent-like states lack long-range order. (c) Experimental data and numerical simulations of the average 2-point correlators squared along the x -axis, $|C_{i,j}^X|^2$, between qubit pairs at distance M for drive duration $t = 10/J$, strength $\Omega = J/2$ and detuning δ from the lattice frequency. (d) The correlation length ξ^X are extracted with an exponential fit following $|C_{i,j}^X|^2 \propto \exp(-M/\xi^X)$ using all measured 2-point correlators as a function of the qubit pair distance. The reported correlation length fit value is accompanied by the standard fit error.

-
- [1] S. Huang, B. Lienhard, G. Calusine, A. Vepsäläinen, J. Braumüller, D. K. Kim, A. J. Melville, B. M. Niedzielski, J. L. Yoder, B. Kannan, T. P. Orlando, S. Gustavsson, and W. D. Oliver, [PRX Quantum](#) **2**, 020306 (2021).
 - [2] F. Yan, S. Gustavsson, A. Kamal, J. Birenbaum, A. P. Sears, D. Hover, T. J. Gudmundsen, D. Rosenberg, G. Samach, S. Weber, J. L. Yoder, T. P. Orlando, J. Clarke, A. J. Kerman, and W. D. Oliver, [Nature Communications](#) **7**, 12964 (2016).
 - [3] C. Macklin, K. O'Brien, D. Hover, M. E. Schwartz, V. Bolkhovskiy, X. Zhang, W. D. Oliver, and I. Siddiqi, [Science](#) **350**, 307 (2015).
 - [4] P. Krantz, M. Kjaergaard, F. Yan, T. P. Orlando, S. Gustavsson, and W. D. Oliver, [Applied Physics Reviews](#) **6**, 021318 (2019).
 - [5] D. Rosenberg, D. Kim, R. Das, D. Yost, S. Gustavsson, D. Hover, P. Krantz, A. Melville, L. Racz, G. O. Samach, S. J. Weber, F. Yan, J. L. Yoder, A. J. Kerman, and W. D. Oliver, [npj Quantum Information](#) **3**, 1 (2017).
 - [6] J. Koch, T. M. Yu, J. Gambetta, A. A. Houck, D. I. Schuster, J. Majer, A. Blais, M. H. Devoret, S. M. Girvin, and R. J. Schoelkopf, [Physical Review A](#) **76**, 042319 (2007).
 - [7] J. Braumüller, A. H. Karamlou, Y. Yanay, B. Kannan, D. Kim, M. Kjaergaard, A. Melville, B. M. Niedzielski, Y. Sung, A. Vepsäläinen, R. Winik, J. L. Yoder, T. P. Orlando, S. Gustavsson, C. Tahan, and W. D. Oliver, [Nature Physics](#) **18**, 172 (2022).
 - [8] A. H. Karamlou, J. Braumüller, Y. Yanay, A. Di Paolo, P. M. Harrington, B. Kannan, D. Kim, M. Kjaergaard, A. Melville, S. Muschinske, B. M. Niedzielski, A. Vepsäläinen, R. Winik, J. L. Yoder, M. Schwartz, C. Tahan, T. P. Orlando, S. Gustavsson, and W. D. Oliver, [npj Quantum Information](#) **8**, 1 (2022).
 - [9] G. J. Dolan, [Appl. Phys. Lett.](#) **31**, 337 (1977).
 - [10] J. Braumüller, L. Ding, A. P. Vepsäläinen, Y. Sung, M. Kjaergaard, T. Menke, R. Winik, D. Kim, B. M. Niedzielski, A. Melville, J. L. Yoder, C. F. Hirjibehedin, T. P. Orlando, S. Gustavsson, and W. D. Oliver, [Phys. Rev. Appl.](#) **13**, 054079 (2020).
 - [11] C. N. Barrett, A. H. Karamlou, S. E. Muschinske, I. T. Rosen, J. Braumüller, R. Das, D. K. Kim, B. M. Niedzielski, M. Schuldt, K. Serniak, M. E. Schwartz, J. L. Yoder, T. P. Orlando, S. Gustavsson, J. A. Grover, and W. D. Oliver, [Phys. Rev. Appl.](#) **20**, 024070 (2023).
 - [12] M. A. Rol, L. Ciorciaro, F. K. Malinowski, B. M. Tarasinski, R. E. Sagastizabal, C. C. Bultink, Y. Salathe, N. Haandbaek, J. Sedivy, and L. DiCarlo, [Applied Physics Letters](#) **116**, 054001 (2020).
 - [13] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos, D. Cournapeau, M. Brucher, M. Perrot, and É. Duchesnay, [Journal of Machine Learning Research](#) **12**, 2825 (2011).
 - [14] M. R. Geller and M. Sun, [Quantum Science and Technology](#) **6**, 025009 (2021).
 - [15] A. Blais, A. L. Grimsmo, S. M. Girvin, and A. Wallraff, [Reviews of Modern Physics](#) **93**, 025005 (2021).
 - [16] J. Bradbury, R. Frostig, P. Hawkins, M. J. Johnson, C. Leary, D. Maclaurin, G. Necoara, A. Paszke, J. VanderPlas, S. Wanderman-Milne, and Q. Zhang, [JAX: composable transformations of Python+NumPy programs](#) (2018).
 - [17] I. Babuschkin, K. Baumli, A. Bell, S. Bhupatiraju, J. Bruce, P. Buchlovsky, D. Budden, T. Cai, A. Clark, I. Danihelka, A. Dedieu, C. Fantacci, J. Godwin, C. Jones, R. Hemsley, T. Hennigan, M. Hessel, S. Hou, S. Kapturowski, T. Keck, I. Kemaev, M. King, M. Kunesch, L. Martens, H. Merzic, V. Mikulik, T. Norman, G. Papamakarios, J. Quan, R. Ring, F. Ruiz, A. Sanchez, R. Schneider, E. Sezener, S. Spencer, S. Srinivasan, W. Stokowiec, L. Wang, G. Zhou, and F. Viola, [The DeepMind JAX Ecosystem](#) (2020).
 - [18] Y. Yanay, J. Braumüller, S. Gustavsson, W. D. Oliver, and C. Tahan, [npj Quantum Information](#) **6**, 1 (2020).
 - [19] D. A. Meyer and N. R. Wallach, [Journal of Mathematical Physics](#) **43**, 4273 (2002).
 - [20] Qiskit contributors, [Qiskit: An open-source framework for quantum computing](#) (2023).
 - [21] E. Pavarini and E. Koch, *Topology, Entanglement, and Strong Correlations*, Schriften Des Forschungszentrums Jülich. Reihe Modeling and Simulation No. 10 (Forschungszentrum Jülich GmbH Zentralbibliothek, Verlag, Jülich, 2020).
 - [22] M. A. Nielsen and I. L. Chuang, *Quantum Computation and Quantum Information: 10th Anniversary Edition*, anniversary edition ed. (Cambridge University Press, Cambridge ; New York, 2011).
 - [23] J. Eisert, M. Cramer, and M. B. Plenio, [Reviews of Modern Physics](#) **82**, 277 (2010).
 - [24] S. Gustavsson, J. Bylander, F. Yan, P. Forn-Díaz, V. Bolkhovskiy, D. Braje, G. Fitch, K. Harrabi, D. Lennon, J. Miloshi, P. Murphy, R. Slattey, S. Spector, B. Turek, T. Weir, P. B. Welander, F. Yoshihara, D. G. Cory, Y. Nakamura, T. P. Orlando, and W. D. Oliver, [Phys. Rev. Lett.](#) **108**, 170503 (2012).
 - [25] F. Yan, S. Gustavsson, J. Bylander, X. Jin, F. Yoshihara, D. G. Cory, Y. Nakamura, T. P. Orlando, and W. D. Oliver, [Nature communications](#) **4**, 2337 (2013).
 - [26] M. Srednicki, [Physical Review E](#) **50**, 888 (1994).