

Supplementary Materials for

A bio-inspired hybrid lifelong learning based on self-selective memristors

Xuemeng Fan^{1,5#}, Guobin Zhang^{1,2,5#}, Zhejia Zhang^{1,5,8}, Xinheng Mei⁴, Zijian Wang^{1,5,8}, Wenjue Zhou^{1,5}, Daying Sun⁴, Lei Deng⁶, Mingkun Xu⁷, Shuai Zhong⁷, Yi Tong⁸, Bin Yu¹, Rong Zhao⁶, Dashan Shang^{3*}, Xiaolei Zhu^{1*}, Qing Wan^{2*}, Yishu Zhang^{1,2,5*}

¹School of Integrated Circuits, Zhejiang University, ZJU-Hangzhou Global Scientific and Technological Innovation Center, Hangzhou 310027, China

²Yongjiang Laboratory, Ningbo, Zhejiang 315202, China

³Key Lab of Fabrication Technologies for Integrated Circuits, Institute of Microelectronics, University of Chinese Academy of Science, 100029, Beijing, China

⁴School of Microelectronics, Nanjing University of Science and Technology, Nanjing, 210094, China

⁵Zhejiang ICsprout Semiconductor Co., Ltd, Hangzhou, Zhejiang, 310027, China

⁶Department of Precision Instruments, Tsinghua University, Beijing, 100084, China

⁷Guangdong Institute of Intelligence Science and Technology, Hengqin, Zhuhai 519031, Guangdong, China

⁸Suzhou Laboratory, Suzhou, Jiangsu, 215123, China

*Email: shangdashan@ime.ac.cn; xl_zhu@zju.edu.cn; qingwan@ylab.ac.cn; zhangyishu@zju.edu.cn

Contributed equally to this work

The PDF file includes:

Note S1. Modeling the Scalability of the Proposed SSMs

Note S2. Mathematical model of the bio-inspired SSM-based model

Note S3. Calibration strategies of conductance errors generated by nonlinearity operations and peripheral circuits

Note S4. Evaluation Metrics

Note S5 Evaluation Protocol and Benchmark Datasets

Note S6. Computation of CPU/GPU performance benchmark

Fig. S1 | Elemental distribution of the SSM device.

Fig.S2 Interface analysis and thin film characterization of the SSM device.

Fig.S3 Endurance and retention of the device.

Fig.S4 Active forgetting and inter-device variability for the 1Kb Array.

Fig.S5 I-V characteristics of 100 SSM devices.

Fig.S6 Voltage programming scheme and scalability.

Fig.S7 Analog conductance characteristics of the SSM.

Fig.S8 SSM mimics synaptic key functions

Fig.S9 Effect of potentiated stimulation on active forgetting.

Fig.S10 Active forgetting responses to diverse stimulation pulses in terms of amplitude, pulse width, and frequency.

Fig.S11.The potentiation/inhibition response under consecutively set pulses and reset pulses.

Fig.S12. Weight changes of potentiation/inhibition modulated by different set/reset voltages.

Fig.S13 The switching endurance and programmability of the developed SSM device.

Fig.S14 Variability of C2C under pulse programming.

Fig.S15 Energy band diagram and characteristic mechanism of the SSM.

Fig.S16 I-V hysteresis curves of single-layer and bilayer GaO_x/InO_x devices.

Fig.S17 Mechanistic fitting.

Fig.S18 Electrical properties of the oxygen vacancy-rich double-layer device.

Fig.S19 Relationship between electrical performance and device size.

Fig. S 20. Schematic diagram of the proposed HLL model architecture based on five SSM processing cores.

Fig.S21 Functional block diagram of the EHLL system.

Fig.S22 Results of MAC operation performed by the SSM array.

Fig.S23 Effect of voltage amplitude on weight relaxation of 1Kb array.

Fig.S24 Active forgetting characteristics of 5 SPCs.

Fig.S25 Statistical analysis of the corresponding relationship between weight states and decay factor.

Fig.S26 Demonstration of the balance between LP and MS under different pulse widths and frequencies.

Fig.S27 Consistency between the modulation of the custom decay factor α and the weight update of SPCs.

Fig.S28 Weight mapping after hybrid training.

Fig.S29 Discrimination errors of three lifelong learning models based on SSM processing cores on three datasets.

Fig.S30 Average test accuracy of the HLL model learned on three datasets.

Fig.S31| Model generalization and robustness.

Other Supplementary Material for this manuscript includes the following:

Movie S1

Note S1. Modeling the Scalability of the Proposed SSMs.

To assess the scalability of the developed SSMs, we exclusively utilize read margin as the core evaluation parameter in this research. Given the need for precise simulation results, SR, NL, and device on/off ratio are all incorporated into the read margin computation framework. Notably, read margin serves as a critical safeguard for information integrity: during read operations on memristor crossbar arrays, it ensures that stored data can be retrieved accurately despite ambient noise or signal perturbations. A larger read margin translates to enhanced robustness and dependability of the crossbar array in data reading, minimizing the risk of erroneous data interpretation. For the calculation of read margin, we adopt the single bit-line pull-up strategy (OBPU), whose specific mechanism is elaborated in the following section:

$$\text{Read Margin} = \frac{\Delta V}{V_{pu}} = R_{pu} \times \left(\frac{1}{R_1 + R_{pu}} - \frac{1}{R_2 + R_{pu}} \right) \quad (\text{S1})$$

$$R_{sneak} = \frac{2R_p}{(n-1)} + \frac{R_n}{(n-1)^2} \quad (\text{S2})$$

$$R_1 = R_{LRS} \times \left(\frac{R_{sneak}}{R_{LRS} + R_{sneak}} \right) \quad (\text{S3})$$

$$R_2 = R_{HRS} \times \left(\frac{R_{sneak}}{R_{HRS} + R_{sneak}} \right) \quad (\text{S4})$$

where ΔV refers to the difference between the input voltages of the selected cell at LRS and HRS. V_{pu} denotes the pull-up voltage applied to the selected BL, and R_{pu} refers to the corresponding pull-up resistance. R_{sneak} represents the equivalent resistance along the sneak current paths within the memory array. R_p and R_n refer to the resistance of semi-selected and fully unselected memory cells, respectively. The parameter n indicates the total number of WLS or BLs in the entire crossbar array. In matrix memories, crosstalk effects reduce read margins and undermine read operation accuracy, with margin degradation worsening as array size increases. The industry consensus holds that the maximum feasible array size for a memory is defined by the point where its read margins fall below 10% under extreme conditions¹.

Note S2. Mathematical model of the bio-inspired SSM-based model

The conductance updating process in SSMs is closely linked to the parameter update mechanism in the hybrid model. In this hybrid model, parameter updates depend not only on the learning requirements of the current task but also on the preservation of memories from previous tasks and the collaborative learning between new and old tasks. This complex update process is implemented through conductance changes in memristors, thereby simulating the learning mechanisms of biological neural systems at the hardware level. The conductance update formula for SSMs is as follows:

$$G(t+1) = G(t) + \Delta G \quad (\text{S5})$$

The variation ΔG integrates the influence of three key components: the gradient information from the current task, the MS regularization term, and the active forgetting term. This ensures that the parameter update process not only adapts to the learning requirements of new tasks but also prevents excessive forgetting of knowledge from previous tasks. In this way, the conductance changes in memristors work in concert with the parameter update mechanism in the hybrid model, enabling efficient management of the learning process. Separately, for

Drosophila-inspired model, the parameter update formula for each learning module L_i is demonstrated as equation:

$$\Delta\theta_{AF} = -\eta\mu_i \left[\nabla_{\theta_i} \mathcal{L}_i + \lambda_{MS} \sum_{j \neq i} (\theta_i - \theta_j) + \lambda_{AF} \alpha (\theta_i - \theta_{i,old}) \right] \quad (S6)$$

In the proposed framework, $\mathcal{L}_i$ denotes the loss function of hardware module S_i , which evaluates the performance on the current task. The terms λ_{MS} and λ_{AF} represent the regularization weights for MS and active forgetting, respectively. The parameter α is a custom decay factor designed to control the forgetting rate of old memories, while $\theta_{i,old}$ corresponds to the parameters of hardware module S_i before learning the new task. Furthermore, we incorporated a neuromodulation factor μ_i to dynamically adjust the learning rate, emulating the neuromodulatory mechanisms found in biological neural systems². Equation S7 translates this parameter-level update to the hardware conductance domain. Combined with the HB-inspired model, the parameter update formula for each learning hardware module is as follows²:

$$\Delta G_{HCL} = (1 - \alpha)\eta\mu_i f_{local}(\xi_i) \left(\mathcal{F}_{i,j}^{new} (G_{i,j} - G_{i,j}^{new}) \right) + \alpha\eta\mathcal{F}_{i,j}^{old} (G_{i,j} - G_{i,j}^{old}) \quad (S7)$$

where $f_{local}(\xi_i)$ denotes a nonlinear neural modulation function, depending on the global signal ξ_i , for regulating the magnitude and direction of synaptic updating. $\mathcal{F}_{i,j}^{new}$ and $\mathcal{F}_{i,j}^{old}$ denotes Fisher information matrix for the new task and old task, respectively[3-4], and $G_{i,j}^{new}$ and $G_{i,j}^{old}$ represents ideal synaptic weights for new and old tasks, respectively. Here, $G_{i,j}^{new}$ is obtained by mapping the updated parameters $\theta_{i,j}^{new} = \theta_i + \Delta\theta_{AF}$ from Equation S6 through the device's parameter-conductance calibration function. This hardware-enabled system captures the characteristics of multi-module collaborative lifelong learning, incorporating three key components: the gradient of the current tasks' loss function, an inter-module parameter difference regularization term, and an active forgetting term. Through SSM conductance programming, this complex parameter update process is physically implemented at the hardware level. A direct correspondence exists between the SSM's conductance change (ΔG) and the parameter update ($\Delta\theta$) [2]. This correspondence enables SSMs to precisely simulate parameter variations in the hybrid model, thereby achieving efficient learning process management in hardware implementation. In summary, the proposed framework operates as a three-level hierarchical pipeline: (i) Equation S6 computes the desired parameter update $\Delta\theta_{AF}$ in weight space, incorporating gradient descent, memory stability, and active forgetting; (ii) Equation S7 translates $\Delta\theta$ into a conductance update ΔG_{HCL} using Fisher information-weighted metaplasticity; (iii) Equations S8–S10 map the target conductance change to device-level programming pulses with calibrated widths and amplitudes. This cascade enables SSMs to precisely simulate parameter variations in the hybrid model at the hardware level.

Note S3. Calibration strategies of conductance errors generated by nonlinearity operations and peripheral circuits

The conductance tuning protocol for our 32×32 SSM crossbar arrays implements a stochastic pulse-width modulation scheme to compensate for intrinsic device variability during lifelong learning tasks. Unlike conventional 1T1R architecture^{5,6}, the absence of selector transistors

necessitates a modified approach where conductance adjustment relies exclusively on the nonlinear voltage-time product across SSM devices, though there is self-rectifying function. The target conductance matrix G_{target} , derived from the lifelong learning model's weight updates, is decomposed into discrete pulse trains according to the equation as below.

$$G_{target} = \sum_{k=1}^K \rho_k \cdot \mathcal{H}(V_{pulse}, t_k) \quad (S8)$$

where $\mathcal{H}(V_{pulse}, t_k)$ represents the memristance modulation function for a pulse of amplitude V_{pulse} and programmable width t_k , with ρ_k as pulse contribution coefficients calibrated via offline characterization. Although the SSM device supports 10 distinct conductance levels, the stochastic pulse-width modulation scheme achieves effective sub-level programming precision. The pulse width t_k in Equation S8 is a continuous variable, allowing the total conductance change to assume a near-continuous range of values. Furthermore, the iterative three-phase programming protocol (baseline profiling-directional pulsing-convergence validation) employs closed-loop error feedback: after each pulse, the relative error $\epsilon_{relative}$ (Equation S11) is evaluated, and subsequent pulse parameters are adjusted accordingly. This iterative refinement progressively reduces quantization errors below the single-level resolution. The 10-level constraint therefore represents the nominal static resolution of a single write operation, not the effective resolution achievable through multi-pulse accumulation and feedback-controlled tuning.

During fine-tuning, each SSM device undergoes iterative conductance adjustment through a three-phase process: (1) Baseline profiling applies a sub-threshold read voltage to measure initial conductance G_{init} while avoiding unintended state drift; (2) Directional pulsing employs bidirectional pulses with durations scaled by the error (Equation S11), where pulse width as equation (S12).

$$\Delta G = |G_{target} - G_{init}| \quad (S9)$$

$$t_{pulse} = \beta \cdot \log(1 + \Delta G / G_{scale}) \quad (S10)$$

with β and G_{scale} as device-specific constants extracted from prior cycle-to-cycle variability measurements; (3) Convergence validation reassesses conductance after each pulse using a moving-average filter to suppress read noise, terminating adjustment when the relative error satisfies:

$$\epsilon_{relative} = \frac{|G_{measure} - G_{target}|}{\max(G_{target}, G_{min})} \quad (S11)$$

where G_{min} prevents division by near-zero conductance. To address peripheral circuit-induced deviations inherent in transistor-less designs, we implement a dual-reference calibration technique. Two dedicated reference columns (non-switchable SSMs with pre-characterized G_{rel_high} and G_{rel_low}) provide real-time conductance benchmarks. During crossbar programming, their output currents I_{rel_high} and I_{rel_low} are sampled to compute a piecewise linear correction factor $\eta(V_{row})$ for each row voltage:

$$\eta(V_{row}) = \frac{G_{rel_{high}} - G_{rel_{low}}}{\frac{I_{rel_{high}}}{V_{row}} - \frac{I_{rel_{low}}}{V_{row}}} \quad (S12)$$

This compensates for voltage drops along unselected SSM paths—a critical consideration in selector-free arrays. The final calibrated conductance for operational columns is then:

$$G_{operational} = \eta(V_{row}) \cdot \frac{I_{col}}{V_{row}} \quad (S13)$$

Note S4. Evaluation Metrics

Following established continual learning protocols, we evaluate model performance using five complementary metrics^{7,8}:

(1) Average Accuracy (ACC):

$$ACC = \frac{1}{M} \sum_{m=1}^M (A_{M,m}) \quad (S14)$$

, where $A_{M,m}$ is the test accuracy on task m after training on all M tasks. ACC measures the model's overall performance across the entire task sequence.

(2) Forward Transfer (FT):

$$FT = \frac{1}{M-1} \sum_{m=2}^M (A_{m-1,m} - A_{baseline}) \quad (S15)$$

, where $A_{m-1,m}$ is the zero-shot accuracy on task m after training on tasks $1 \dots m-1$, and $A_{baseline}$ is the accuracy of a randomly initialized model trained exclusively on task m . FT measures zero-shot knowledge transfer from prior tasks to a new task, before any training on that task. Positive FT indicates that previously learned representations directly benefit the recognition of new classes without additional training. On R-CIFAR-10/100, positive FT arises because CIFAR-10 pre-training provides transferable low-level visual features that enable above-chance zero-shot classification of CIFAR-100 subclasses, outperforming a randomly initialized model trained on limited per-task data.

(3) Backward Transfer / Memory Stability (BWT / MS):

$$BWT = \frac{1}{M-1} \sum_{m=1}^{M-1} (A_{M,m} - A_{m,m}) \quad (S16)$$

, where $A_{m,m}$ is the accuracy on task m immediately after learning it, and $A_{M,m}$ is the final accuracy after all M tasks. Negative BWT indicates catastrophic forgetting; $BWT \approx 0$ indicates perfect retention. We use BWT and MS interchangeably; MS emphasizes the biological interpretation (memory stability).

(4) Plasticity (Ω_{new}):

Following Kemker et al., Ω_{new} is defined as the area under the per-task learning curve during the initial training phase of each new task, normalized by the number of training iterations. $\Omega_{\text{new}} \in [0, 1]$; higher values indicate faster, more effective acquisition of new knowledge. This metric directly captures the model’s learning plasticity, which is its ability to adapt to novel information.

(5) FT-BWT Ratio:

The ratio $\text{FT} / |\text{BWT}|$ (for $\text{BWT} < 0$) or $\text{FT} / (|\text{BWT}| + \epsilon)$ characterizes the system’s dynamic equilibrium between forward knowledge transfer and memory retention, analogous to the biological stability-plasticity balance.

Note S5 Evaluation Protocol and Benchmark Datasets

All experiments adopt the task-incremental learning (Task-IL) paradigm⁹. In this setting, the dataset is partitioned into M sequentially presented tasks with disjoint class sets. Task identity is provided to the model at both training and test time via a task-specific output head. After training on each task m ($m = 1, \dots, M$), the model is evaluated on the test sets of all tasks seen so far ($1, \dots, m$). No data from previous tasks is stored or replayed during training. Model selection and hyperparameter tuning for each task are conducted on a separate held-out validation split[9].

The construction logic of the datasets we adopt, including R-CIFAR-100, S-CIFAR-100, and R-CIFAR-10/100, reflects a fine-grained control of the key variables in the knowledge migration process: R-CIFAR-100 simulates extreme scenarios without significant semantic associations between tasks by randomly dividing the categories, which is used to test the model’s ability to adapt to dramatic changes in the distribution of the R-CIFAR-100 simulates extreme scenarios with no significant semantic associations between tasks by randomly dividing categories to test the model’s adaptability under drastic changes in distribution; S-CIFAR-100 examines the model’s ability to capture hierarchical knowledge structures by dividing tasks based on semantic super-categories, reflecting the typical real-world scenarios with partial associations between tasks; and R-CIFAR-10/100 constructs progressive learning sequences across the scale of the dataset by mixing the CIFAR-10 and CIFAR-100 tasks for evaluating scalability of the method in heterogeneous task streams. This multidimensional benchmark design not only covers the continuous knowledge transfer spectrum from negative to positive migration, but also reveals the dynamic regulation laws of neural adaptive mechanisms under different migration intensities by controlling variables such as input size, sample size, and task similarity, providing an interpretable empirical basis for theoretical analysis.

Note S6. Computation of CPU/GPU performance benchmark

We develop a heterogeneous computing profiling framework to characterize data movement bottlenecks between parallel accelerators and host processors during deep neural network inference. The instrumentation platform integrates NVIDIA Nsight Systems 2023.3 and Intel VTune Profiler 2023.2 analytical tools, operating under a Windows 10 Professional environment. To achieve comprehensive system observability, the monitoring architecture

combines kernel-level profilers with the Windows Performance Monitor subsystem, enabling correlated analysis across hardware abstraction layers. The NVIDIA toolchain specifically traces CUDA memory transactions and execution timelines of parallel workloads, whereas the Intel analyzer quantifies PCI Express interconnect utilization patterns and processor cache coherence behavior. Platform-specific optimizations involve activating hardware-accelerated GPU task scheduling to reduce timing jitter, accompanied by a dedicated synchronization service that orchestrates measurement tools while minimizing OS-induced profiling artifacts. Within the framework, accelerator energy dissipation is modeled according to equation (S17)¹⁰.

$$E = (P_{exec} - P_{idle}) \times \Delta t_1 \quad (S17)$$

where P_{exec} represents the active power consumption during program execution, P_{idle} denotes the baseline system power in idle state, and Δt_1 corresponds to the measured system latency. All lifelong learning models are employed through parallel computing architectures, utilizing NVIDIA CUDA version 12.1 for optimized tensor operations. Power characterization is performed with high temporal resolution, capturing both computational unit and memory hierarchy energy usage across processing elements. For the parallel accelerator, comprehensive power telemetry was acquired at 100ms sampling intervals through NVIDIA's hardware monitoring interface (nvidia-smi), incorporating both streaming multiprocessor activity and memory subsystem dissipation metrics. Complementary host processor energy profiling is conducted using Intel's proprietary power analysis toolkit (version 3.6), with strict synchronization to computational phases to ensure temporal accuracy. The total energy expenditure was quantified through numerical integration of instantaneous power measurements over the complete execution timeline, as expressed by the energy formulation presented as equation (S18)¹⁰.

$$E = TDP \times CPU_{utilization} \times \Delta t_2 \quad (S18)$$

where TDP represents the processor's thermal design power (in watts), $CPU_{utilization}$ denotes the measured CPU usage percentage during process execution (ranging from 0 to 1), and Δt_2 corresponds to the system latency.

The 3.42× speedup and 91× energy efficiency figures are projections derived from measured hardware primitives:

Per-MAC energy: $E_{op} = 12.7$ fJ (measured: $V_{read} = 0.3$ V, $I_{column} = 84.7$ nA average, $t_{read} = 50$ ns per VMM cycle over 1024 cells)

Per-VMM latency: $t_{op} = 50$ ns (measured with FPGA timer)

Projected system energy for a CIFAR forward pass: $E_{system} = N_{tiles} \times 1024 \times E_{op} + E_{peripheral}$

Baseline GPU energy/latency: measured on NVIDIA RTX 3090 using nvidia-smi power sampling
Ideal projection (zero tiling overhead): 91× energy reduction, 3.42× speedup. With estimated tiling overhead (ADC conversion + FPGA accumulation, ~200 ns per tile): 78× energy reduction, 2.8× speedup.

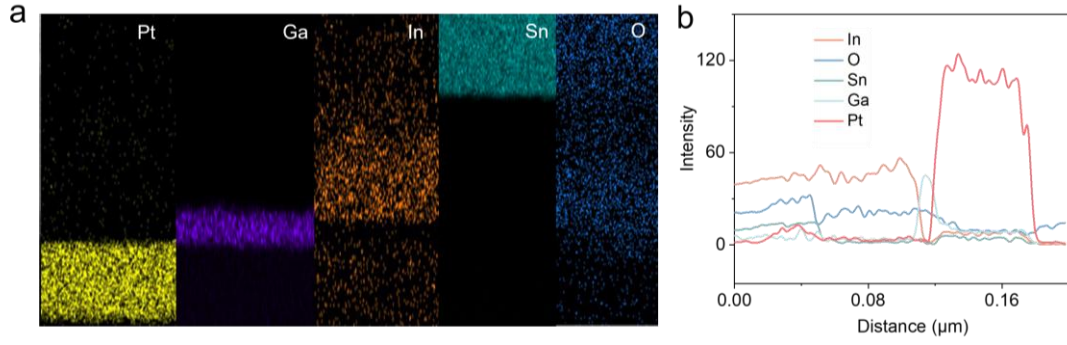


Fig. S1 | Elemental distribution of the SSM device. (A) EDS mapping of the SSM device, displaying the elemental distribution. **(B)** EDS line scanning result of the SSM device.

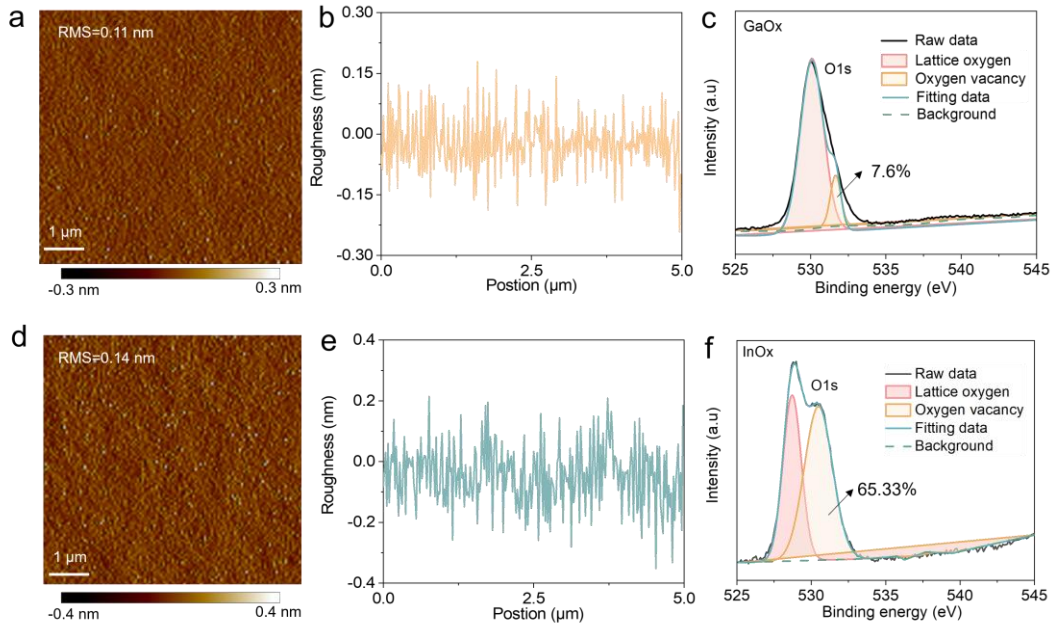


Fig.S2 Interface analysis and thin film characterization of the SSM device. (A) Surface AFM image of the GaO_x film fabricated at a sputtering power of 50 W. Root mean square roughness (RMS). **(B)** Surface profile curve obtained from the scanning line, illustrating the surface roughness characteristics. **(C)** XPS spectrum of O in the GaO_x film. **(D)** Surface AFM image of the InO_x film fabricated at a sputtering power of 50 W. **(E)** Surface profile curve obtained from the scanning line, illustrating the surface roughness characteristics. **(F)** XPS spectrum of O in the InO_x film.

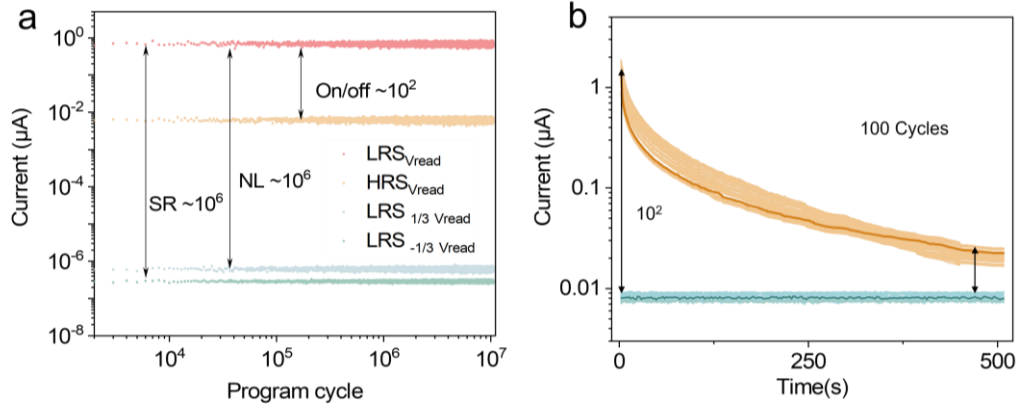


Fig.S3 Endurance and retention of the device. (A) Endurance performance of the SSM device measured over 10^6 consecutive switching cycles; (B) The retention behavior of the SSM, demonstrating that the device spontaneously transitions from the LRS back to the HRS after a retention period of 500 s.

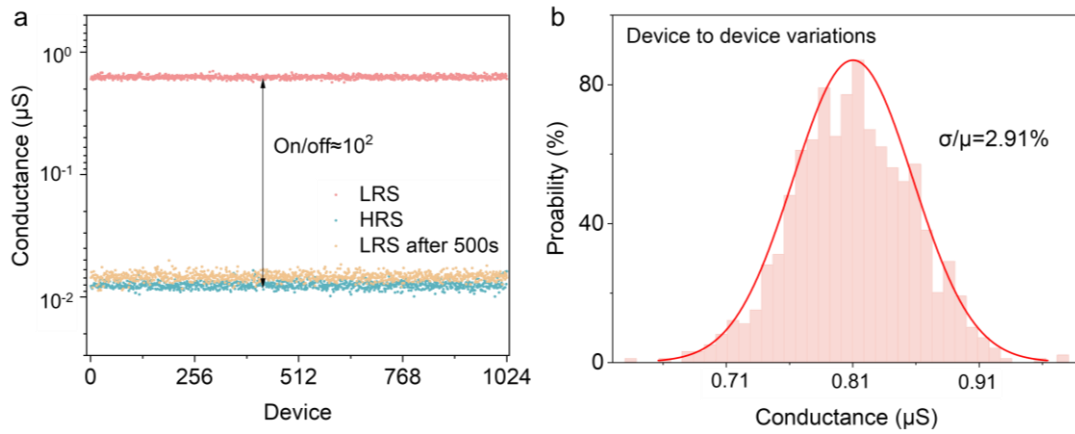
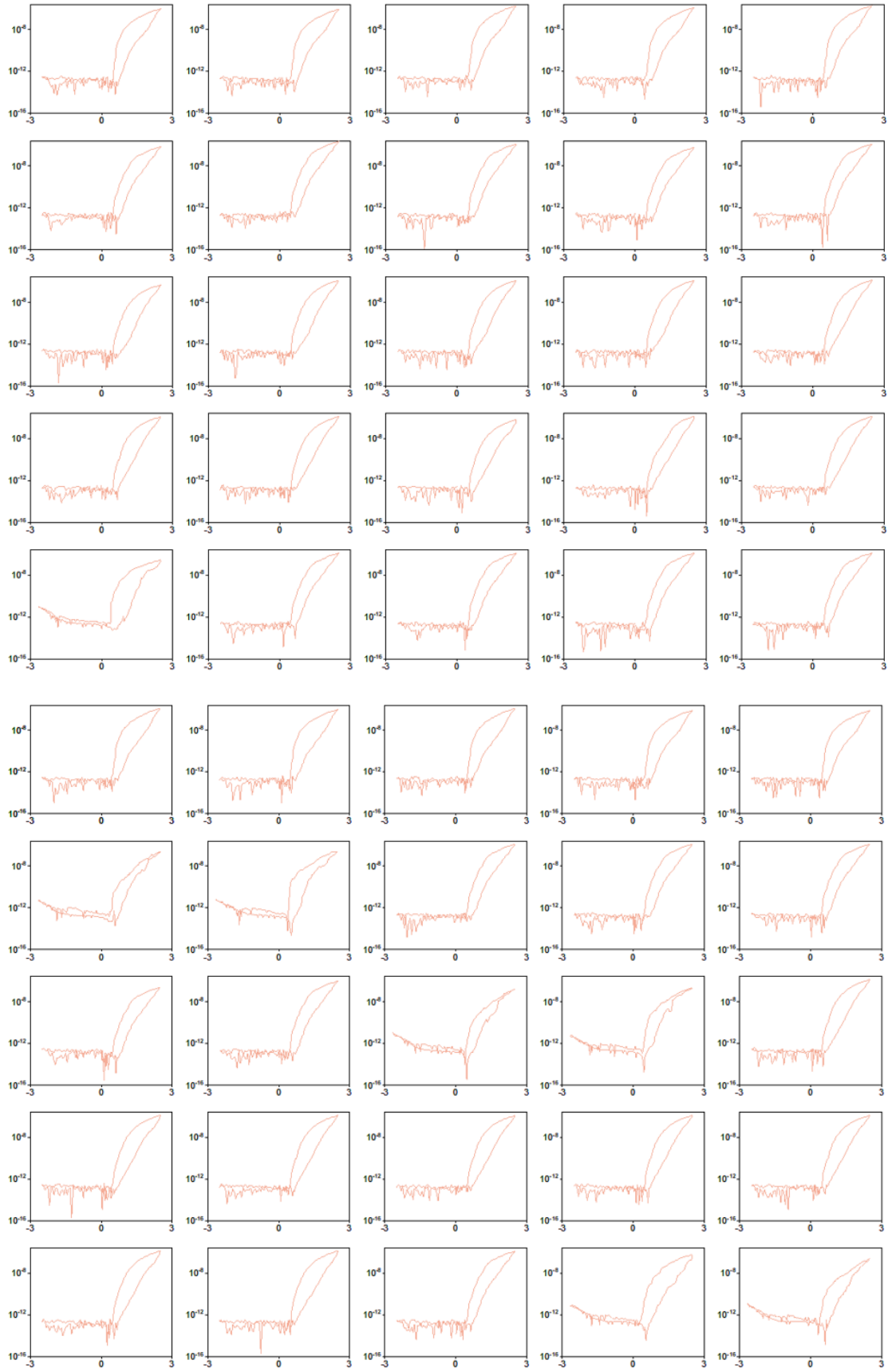


Fig.S4 Active forgetting and inter-device variability for the 1Kb Array. (A) LRS, HRS, and LRS after 500 s of retention for 1024 devices in a 32×32 array. (B) D2D LRS variation (statistics: 1024 devices in 32×32 array).



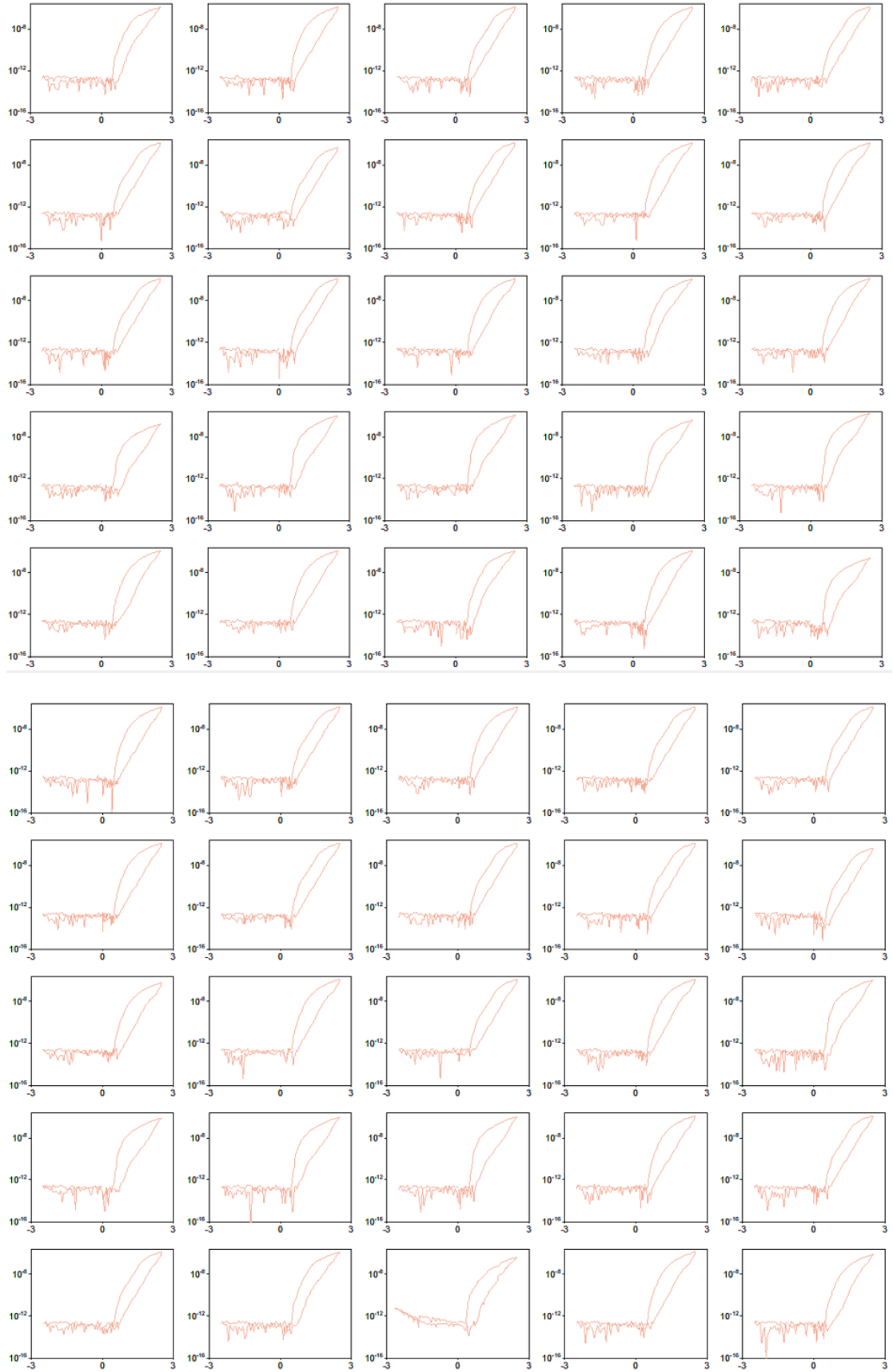


Fig.S5 I-V characteristics of 100 SSM devices. I-V curves of 100 randomly selected SSM units in a 32×32 crossbar array, demonstrating minimal D2D variations.

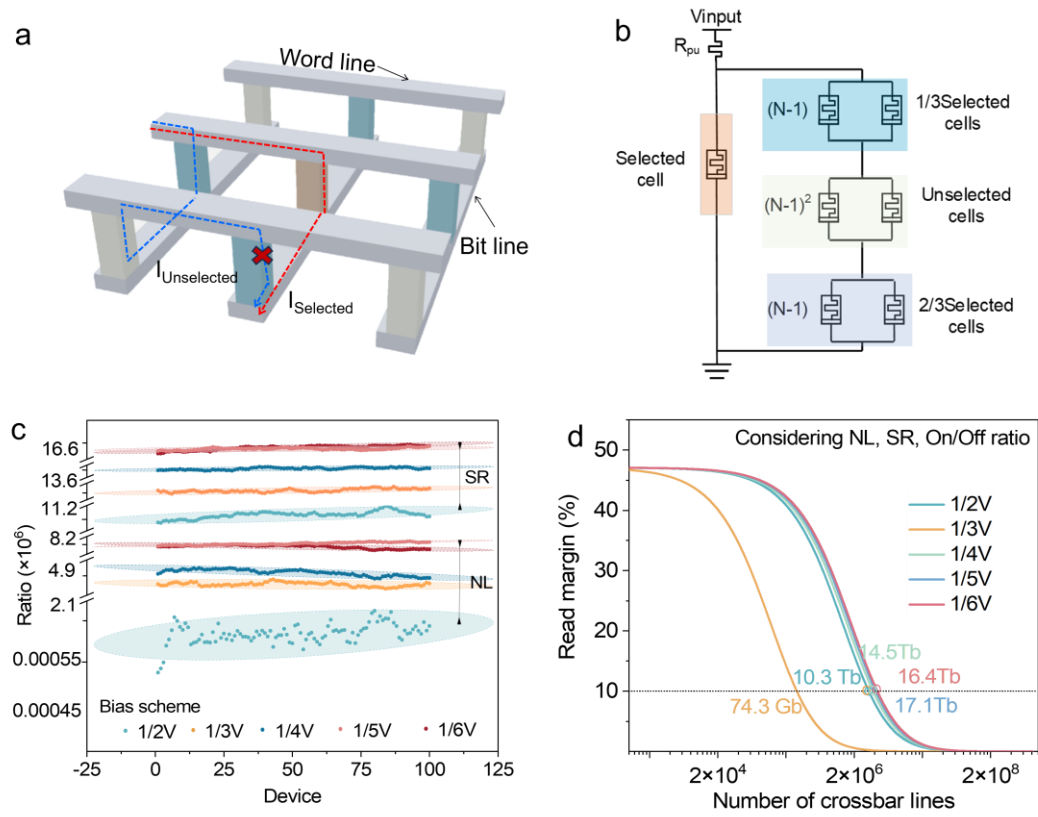


Fig.S6 Voltage programming scheme and scalability. (A) Schematic diagram of sneak path currents in the crossbar array without self-rectifying. (B) Equivalent circuit diagram of one crossbar array with sneak paths. (C) Statistical analysis of NL and SR for 100 devices in the array, obtained with different programming bias schemes. (Dashed lines represent the 95% confidence intervals). (D) The scalability simulation of the SSM, considering NL, SR and On/off ratio, under different programming schemes. Terabit-scale storage density provides relatively excellent support for the large-scale deployment of memristor hardware.

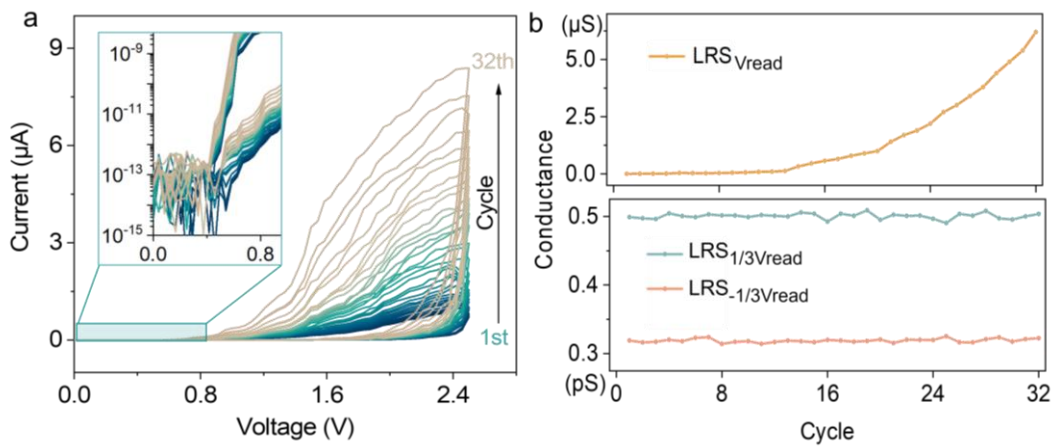


Fig.S7 Analog conductance characteristics of the SSM. (A) I-V characteristic curves under 32 cycles of voltage scanning (0 V-2.5 V-0 V), (B) corresponding resistance state. During the

analog weight update process, the stable NL ($LRS_{vread}/LRS_{1/3Vread}$) and SR ($LRS_{vread}/LRS_{-1/3Vread}$) of the device endow it with the potential to achieve high-precision writing and reading operations.

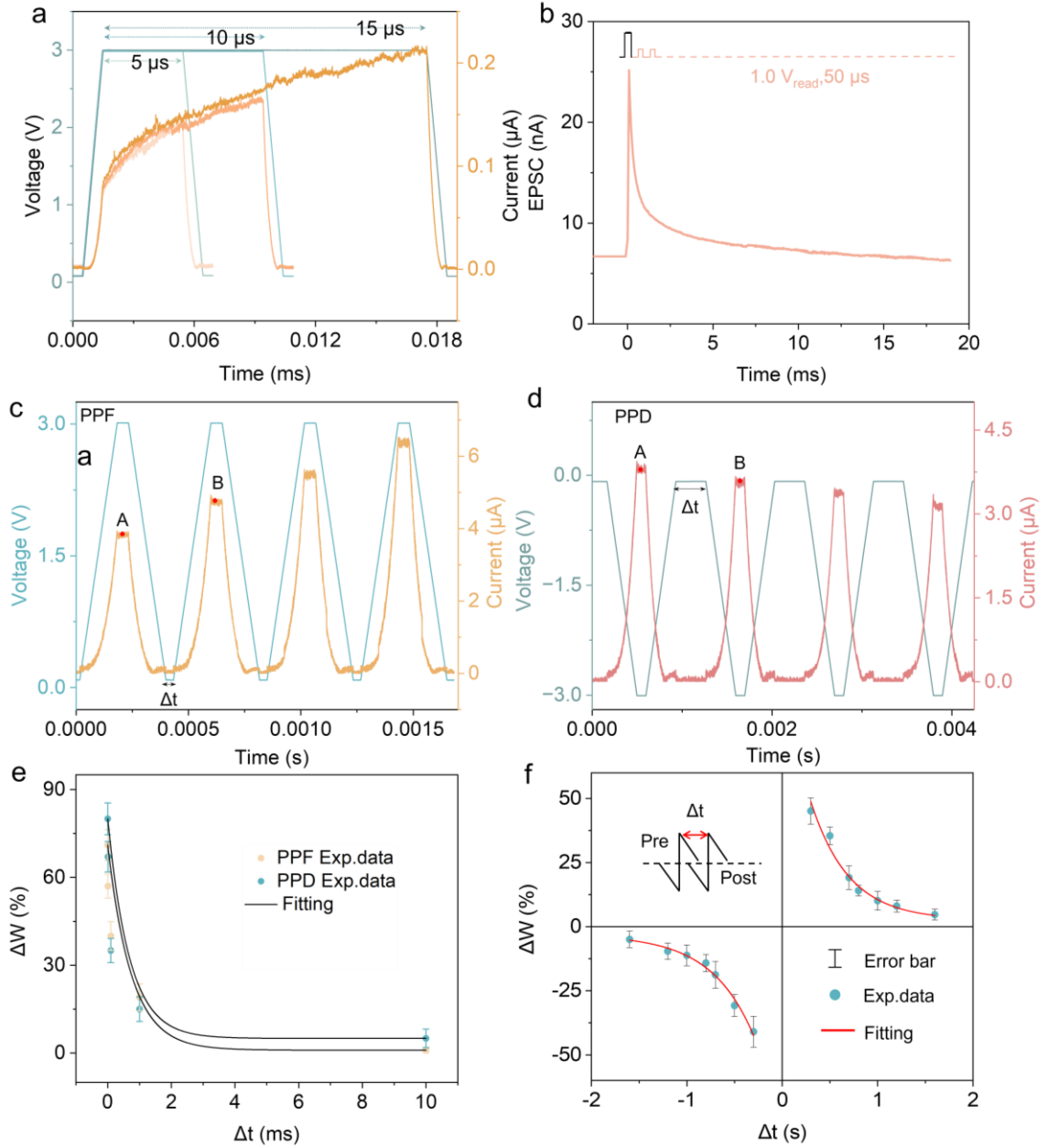


Fig.S8 SSM mimics synaptic key functions. (A) A single pulse stimulates the SSM to generate a current with evident time dependence, simulating the activation of presynaptic neurons. **(B)** the EPSC of the SSM mimics the response of the postsynaptic membrane to presynaptic neuron stimulation. This entire process verifies that the SSM possesses the fundamental information transmission function inherent to synapses. Presynaptic potentiation (3 V,15 μs). **(C)**PPF and **(D)** PPD curves of the SSM device with different pulse intervals. **(E)** The variation of PPF and PPD with the paired voltage pulse interval. We repeated the tests in c and d for 100 times, statistically analyzed the obtained ΔW , and fitting results revealed that the PPF function

exhibited an exponential decay in weight change as the paired-pulse time interval (Δt) increased, which aligns with the characteristic forgetting behavior of biological synapses. $\Delta W = (I_B - I_A) / I_A$. The fitting curve is fitted by the equation, $PPF = C_1 e^{-t/\tau_1} + C_2 e^{-t/\tau_2}$, Where C_1 and C_2 are the initial fitting parameters and τ_1 and τ_2 represent the relaxation time(1). **(F)** Demonstration of STDP in SSM devices. The fitting curve is fitted by the equation(1),

$$\Delta W = \begin{cases} A1 e^{-|\Delta t|/\tau_1}, & \Delta t > 0 \\ A2 e^{-|\Delta t|/\tau_2}, & \Delta t < 0 \end{cases}$$

Where $A1, A2$ is the scaling factor, τ_1 and τ_2 are the time constants of $\pm \Delta t$.

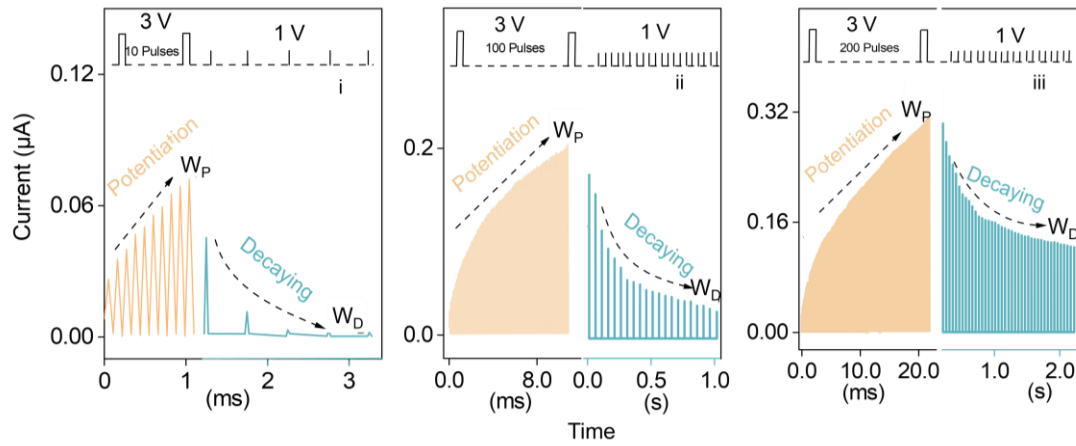


Fig.S9 Effect of potentiated stimulation on active forgetting. Potentiation is induced via 10, 100, and 200 discrete programming pulses (3.0 V, 50 μ s), while subsequent memory decay kinetics are characterized using repetitive read pulses (1 V, 10 μ s).

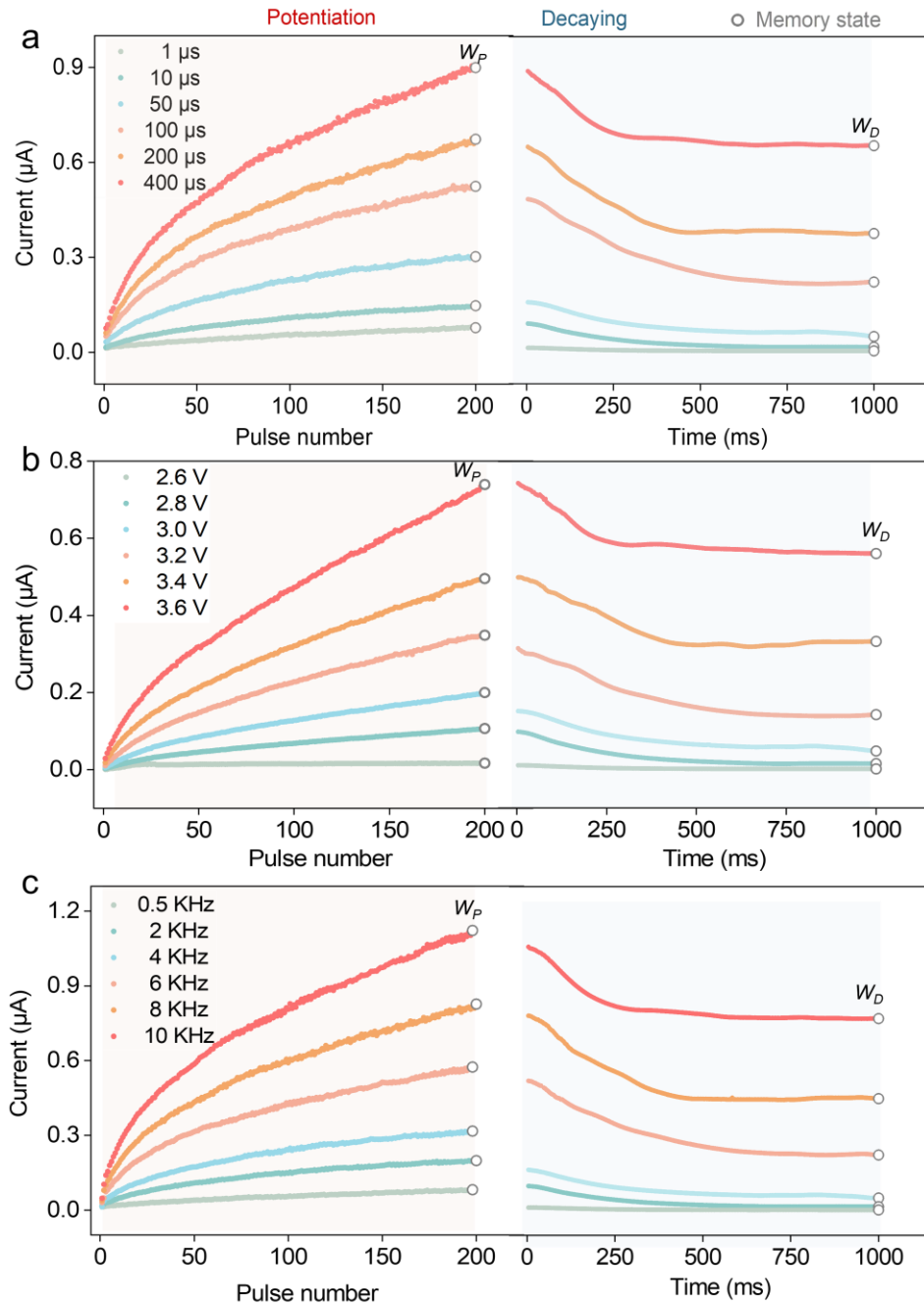


Fig.S10 Active forgetting responses to diverse stimulation pulses in terms of amplitude, pulse width, and frequency. (A) Weight potentiation responses to 200 consecutive set pulses with varying pulse widths (fixed amplitude: 3.0 V). **(B)** eight potentiation responses to 200 consecutive set pulses with varying amplitudes (fixed width: 50 μs). **(C)**Weight potentiation responses to 200 consecutive set pulses with varying frequencies (fixed amplitude: 3.0 V; fixed width: 50 μs). A read pulse (1 V, 10 μs) was used to record the conductance and decay conductance. The decay factor $\alpha=(W_P-W_D)/W_P$ is statistically presented in **Fig. 2F**.

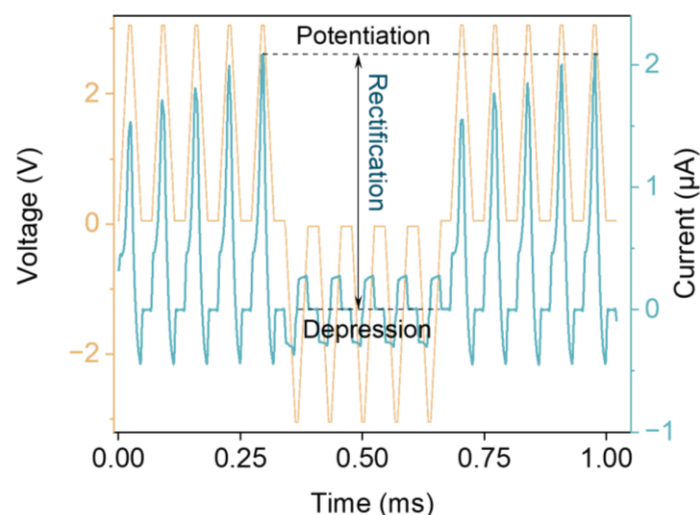


Fig.S11.The potentiation/inhibition response under consecutively set pulses and reset pulses. A low depression current level is sustained by the rectification effect, and the charging current induced during pulse programming gives rise to the observed fluctuating curve.

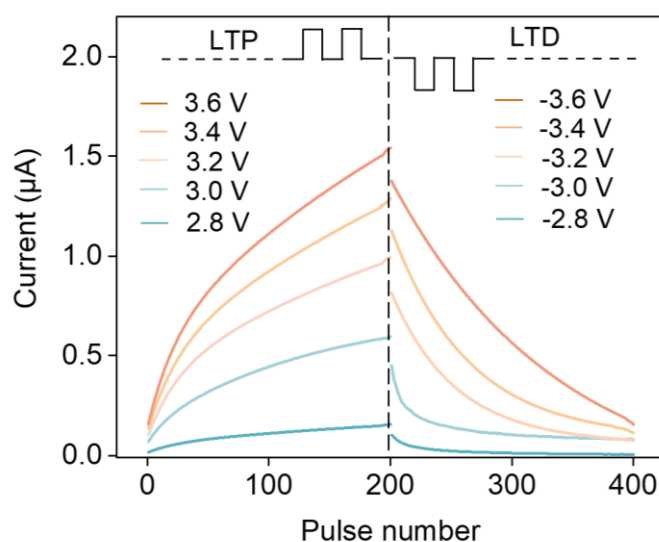


Fig.S12. Weight changes of potentiation/inhibition modulated by different set/reset voltages. A fixed positive read pulse (1 V, 10 μ s) was employed for the measurement. The potentiation and inhibition response currents, which were modulated by the set/reset pulse sequence, were recorded under this read condition.

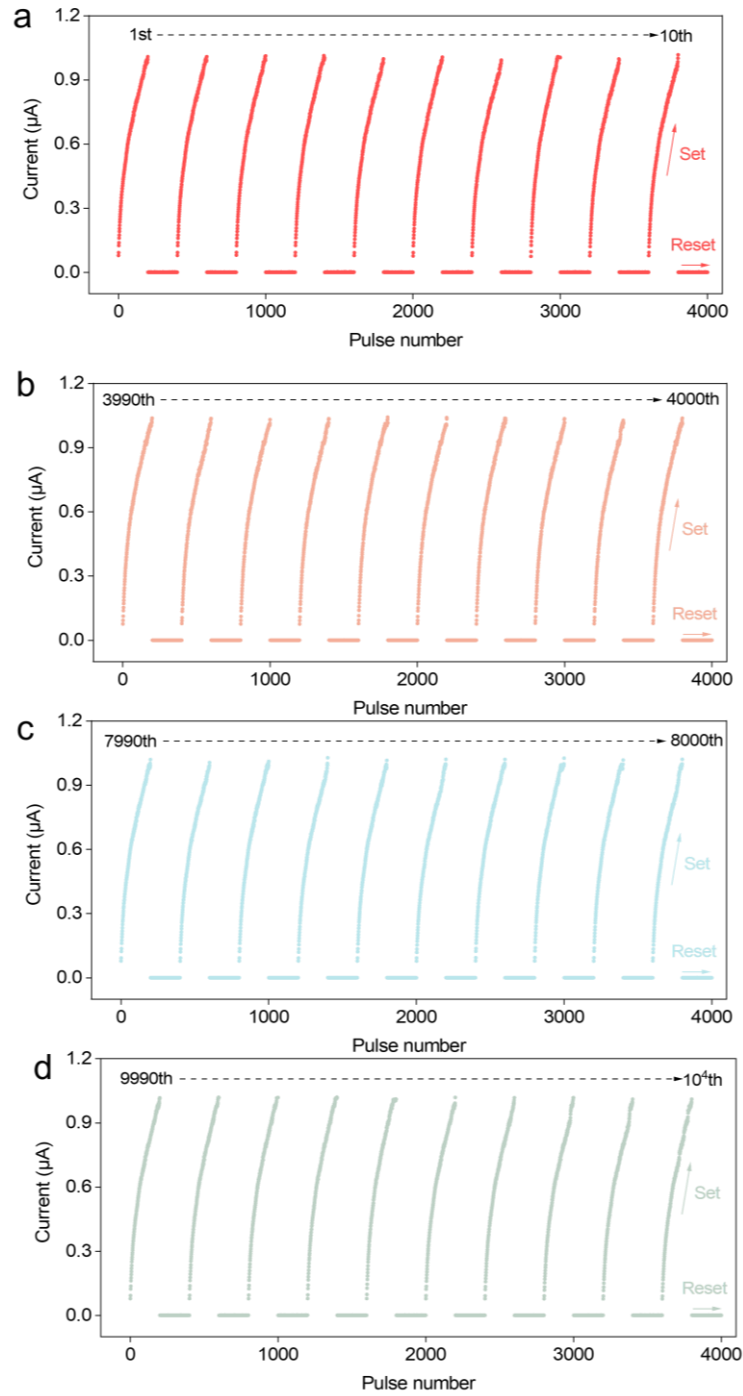


Fig.S13 The switching endurance and programmability of the developed SSM device.

Each test cycle comprised 200 consecutive set pulses (3.0 V, 50 μ s) followed by 200 consecutive reset pulses (-3.0 V, 50 μ s), with the entire sequence repeated for 10^4 cycles. Readout was performed using set and reset read pulses of ± 1 V and 10 μ s, respectively. The figure presents the measured current over the preceding 10 cycles corresponding to the 10th, 4000th, 8000th, and 10^4 th cycles.

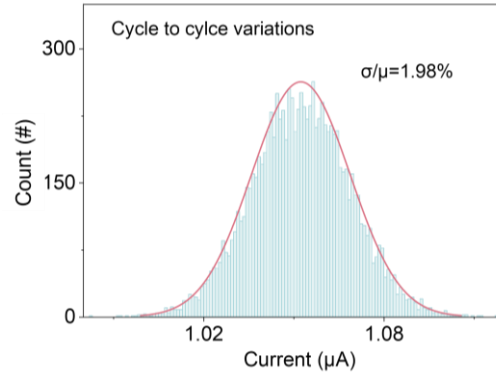


Fig.S14 Variability of C2C under pulse programming. C2C LRS variation is derived from conductance statistics of 200 pulses over 10^4 cycles for one device.

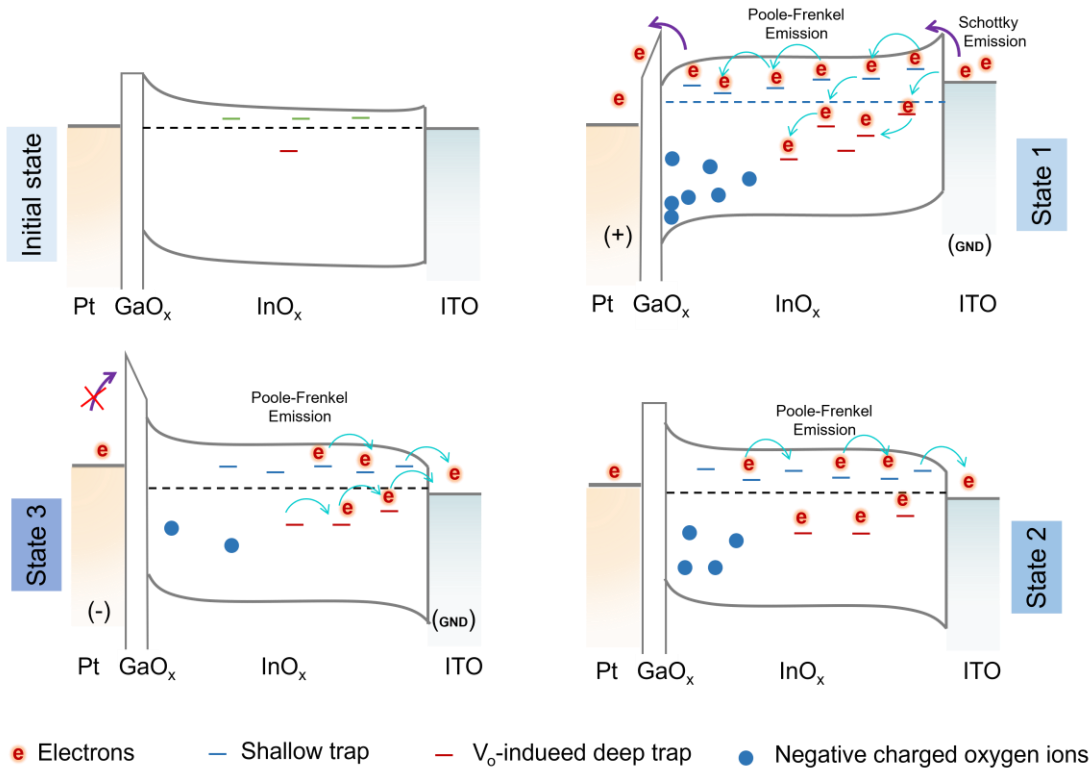


Fig.S15 Energy band diagram and characteristic mechanism of the SSM In the initial state, oxygen vacancies (V_0) in InO_x are sparsely distributed, limiting carrier mobility. Concurrently, the $\text{GaO}_x/\text{InO}_x$ band-offset heterojunction barrier and InO_x/ITO work-function-mismatch Schottky barrier form a dual barrier that blocks electron transport, confining the device to a high-resistance state (HRS). Under positive bias (State 1), the dual barrier sustains HRS with nonlinear I–V response at low voltages (~ 0.5 V). Higher bias drives oxygen ions to $\text{GaO}_x/\text{InO}_x$ and V_0 to InO_x/ITO , restructuring defects. This reduces the InO_x/ITO Schottky barrier for electron injection: some electrons conduct via rapid shallow-trap capture/detrapping, others are slowly trapped in deep levels, while high-energy electrons thermally tunnel to Pt, switching to low-resistance state (LRS). After positive bias removal (State 2), shallow-trap electrons

thermally detrapp and return to ITO (short-term memory decay), whereas deep-trap electrons persist. Residual InO_x carriers prevent HRS reversion, retaining short-term memory weight. Under negative bias (State 3), the reverse field only detraps residual InO_x electrons to ITO. Dense insulating GaO_x forms a high Pt/ GaO_x Schottky barrier, causing low leakage and self-rectification (gradual long-term memory decay). Multiple negative pulses restore initial HRS, enabling short/long-term memory switching.

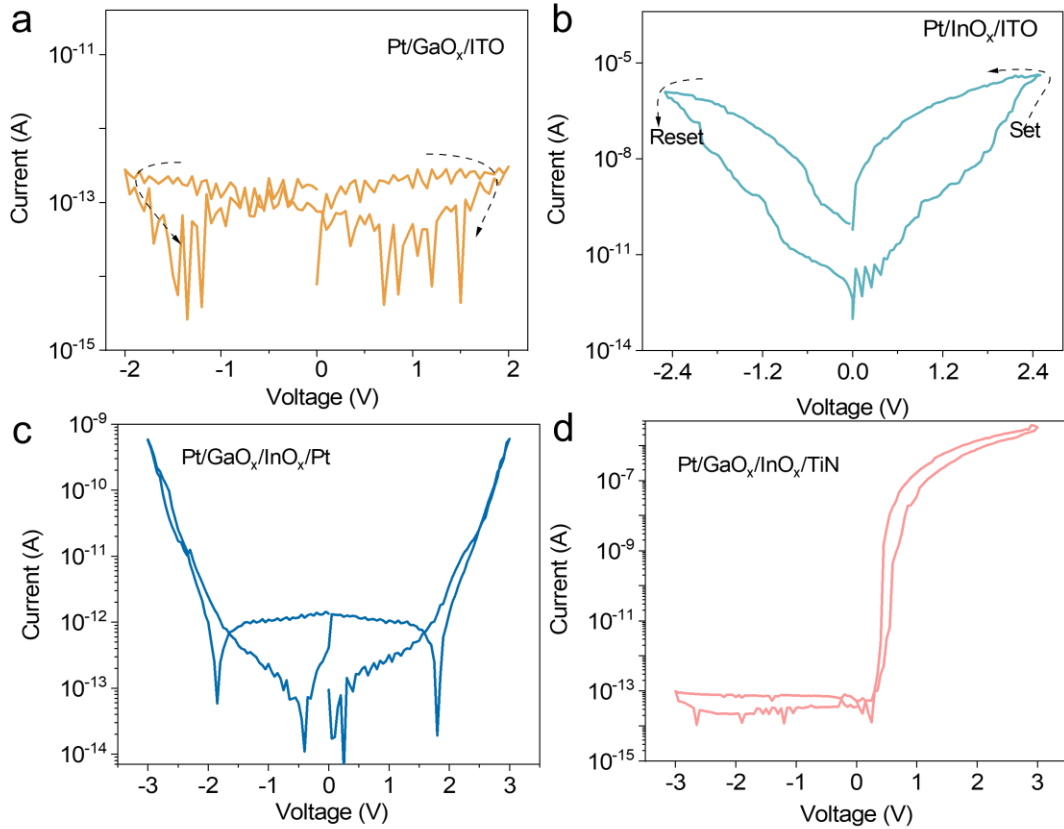


Fig.S16 I-V hysteresis curves of single-layer and bilayer $\text{GaO}_x/\text{InO}_x$ devices. (A) DC I-V curve of Pt/ GaO_x /ITO device and (B) Pt/ InO_x /ITO device. (C), (D) Comparative I-V characteristics of bilayer Pt/ GaO_x /InO_x heterostructures with different top electrodes (Pt and TiN). All film stacking parameters and testing conditions are kept consistent except the top electrode material.

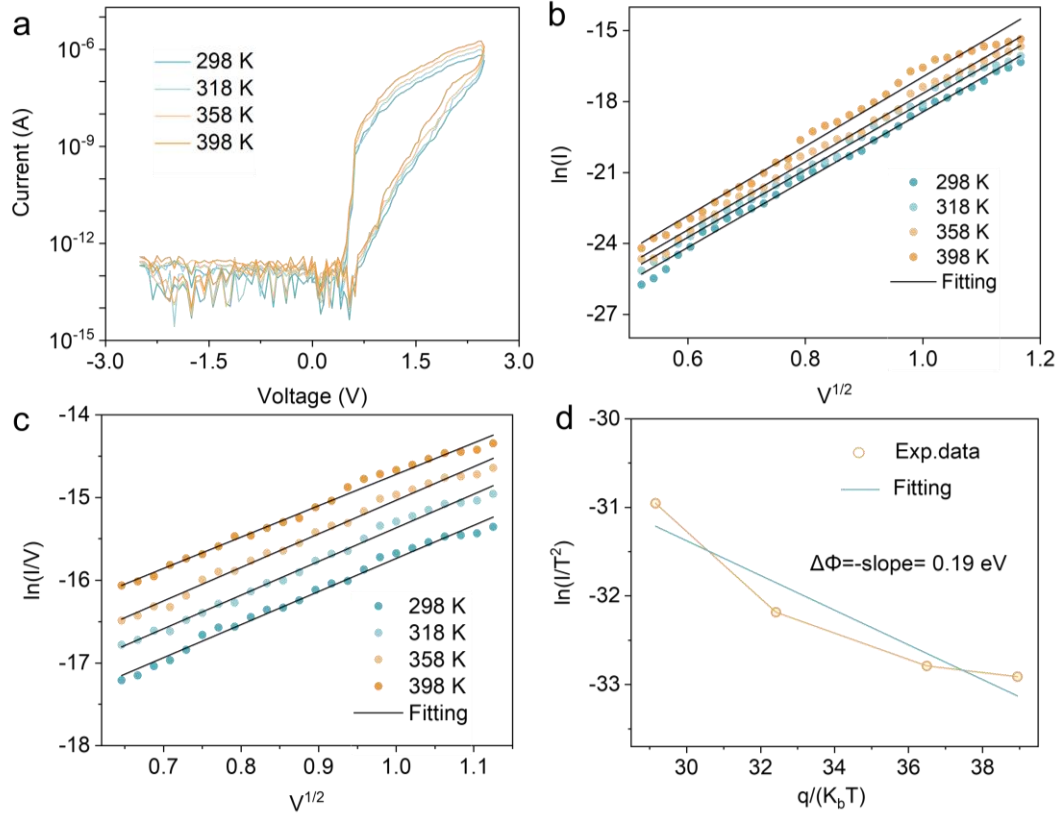


Fig.S17 Mechanistic fitting. (A) DC I-V curves of SSMs at different temperatures. (B) Linear fitting of Schottky emission at different temperatures. (C) Linear fitting of P-F emission at different temperatures. (D) Scatterplot of $\ln(I/T^2)$ versus $1/kT$ with linear fit. The Schottky barrier is linearly associated with the linear fit slope, and Schottky emission essentially involves thermally activated electrons injecting from an energy barrier into the oxide conduction band. This process is described by the equation (S120):

$$J = \beta \mu T^{3/2} E \left(\frac{m^*}{m_0} \right)^{3/2} \exp \left[\frac{-q \left(\phi - \sqrt{\frac{qE}{4\epsilon\pi}} \right)}{kT} \right] \quad (S20)$$

where β represents a constant, μ is the electron drift mobility, m_0 is electron mass, m^* is the electron effective mass, k is Boltzmann's constant, T is the absolute temperature, E is the applied electric field across the oxide, ϕ is the junction barrier height and ϵ is the oxide permittivity. The core principle of P-F emission lies in the excitation of electrons into the oxide's conduction band. The mathematical expression for this phenomenon is given by equation:

$$J_{PF} = q \mu N E \exp \left[\frac{-q \left(\phi_T - \sqrt{\frac{qE}{\pi\epsilon}} \right)}{kT} \right] \quad (S21)$$

Wherein μ is electron drift mobility, N is conduction band state density, E is the oxide-crossing

applied electric field, N is trapped potential well depth, k is Boltzmann's constant, and T is absolute temperature. At higher temperatures, thermal excitation enables carriers to bypass the energy restriction in the high-barrier region, which boosts the current contribution ratio of this region and thereby results in a global increase in the Schottky barrier.

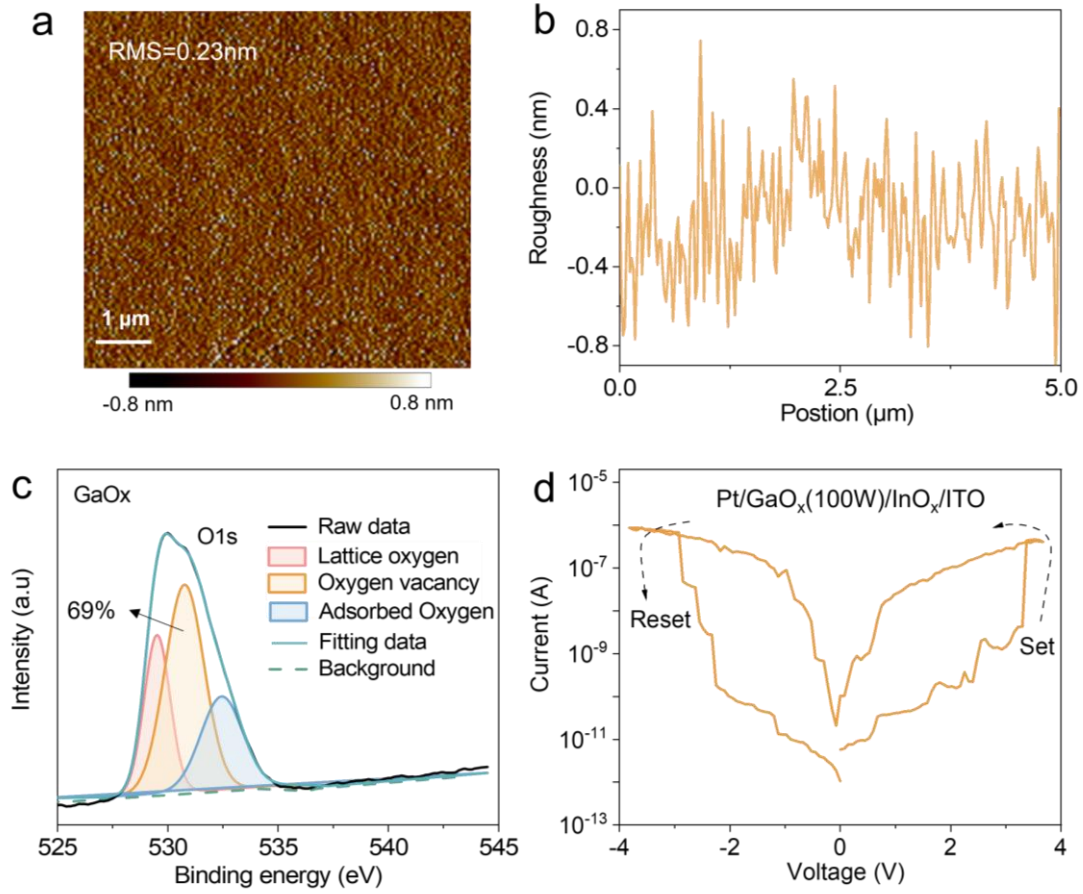


Fig.S18 Electrical properties of the oxygen vacancy-rich double-layer device. (A) Surface AFM image of the GaO_x film fabricated at a sputtering power of 100 W. **(B)** Surface profile curve obtained from the scanning line, illustrating the surface roughness characteristics. **(C)** XPS spectrum of O in the GaO_x film. **(D)** DC I-V curve of the $\text{Pt/GaO}_x/\text{InO}_x/\text{ITO}$ device, with the GaO_x film fabricated at a sputtering power of 100 W.

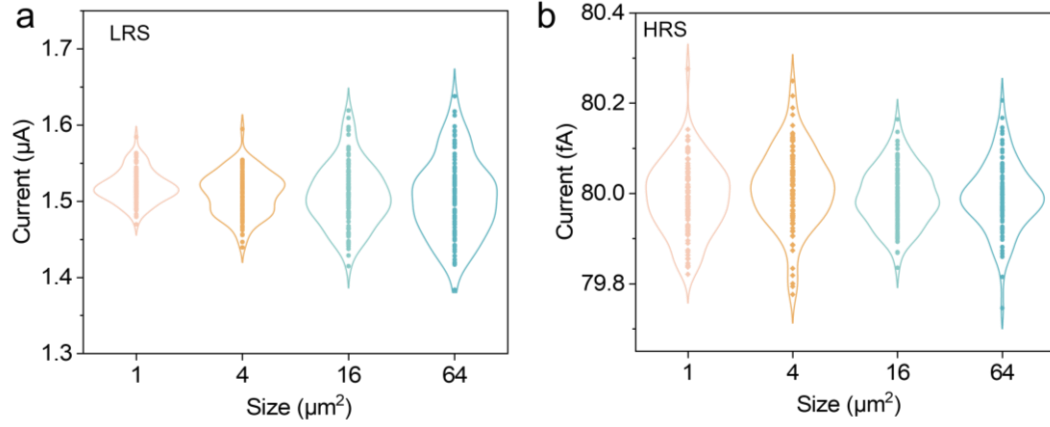


Fig.S19 Relationship between electrical performance and device size. Statistical diagram of the correspondence between (A) LRS (set state) and (B) HRS (reset state) for different device sizes.

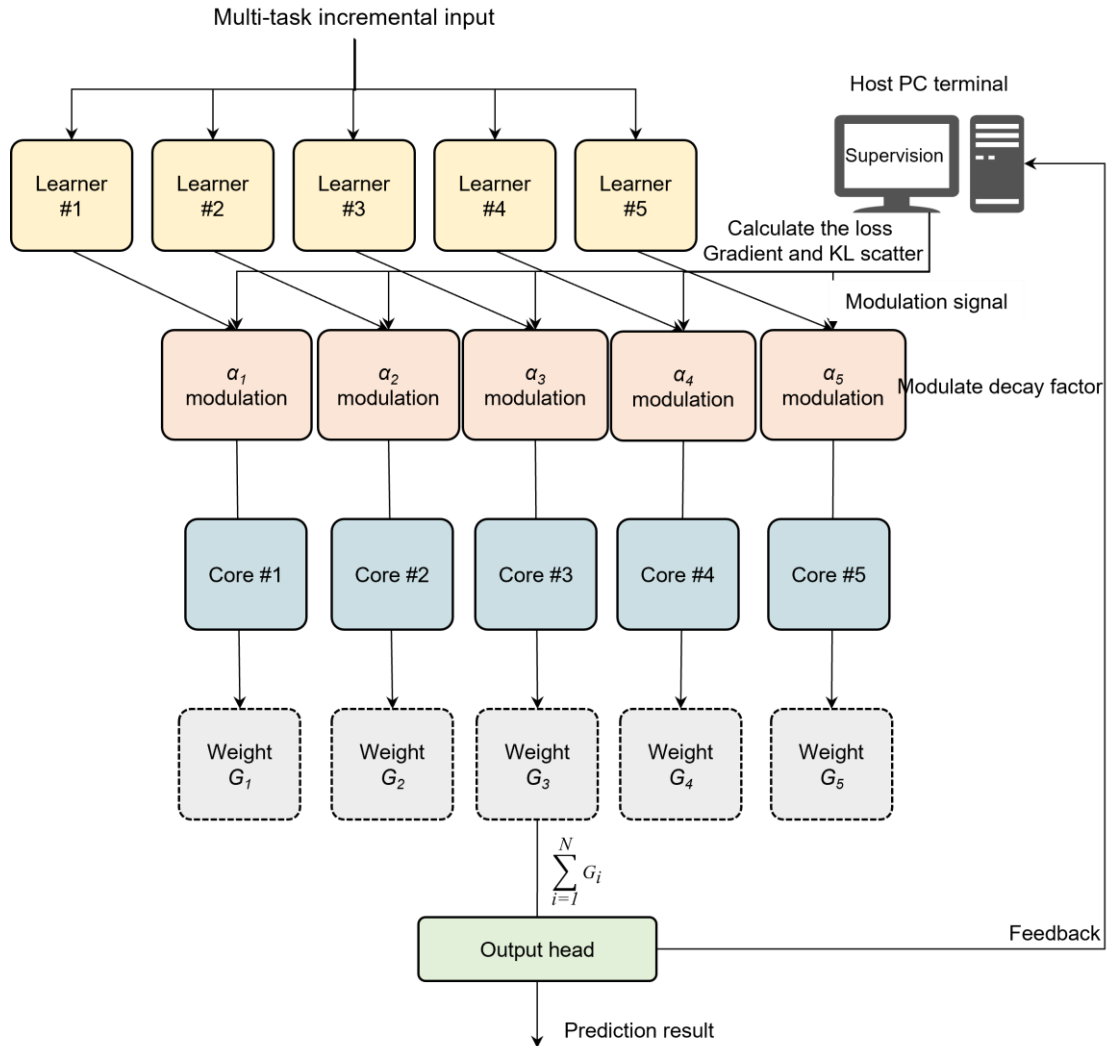


Fig. S 20. Schematic diagram of the proposed HLL model architecture based on five SSM processing cores. This architectural schematic demonstrates a bionic lifelong learning system

based on the principle of neural computation in the γ -region compartment of the *Drosophila* mushroom body, using five SSM cores with the same structure to process the input signals in parallel. Each SSM core implements analog storage and matrix operations of synaptic weights by differential conductance pairs (G^+/G^-), whose conductance values follow the AF characteristic. A voltage-regulated α modulator is integrated inside the core, which receives the analog signal V_a generated by the adaptive controller and realizes the dynamic control of the decay factor α by adjusting the leakage current strength of the SSM unit. The outputs of each core are weighted and summed by the programmable SSM weights $G_1 \sim G_5$, which use polymorphic conductance to achieve multilevel accuracy storage, and their update mechanism is directly written to the SSM unit through pulse amplitude modulation. The integrated signal inputs share the output layer, which consists of a fully-SSM network that utilizes threshold shift characteristics to achieve a ReLU-like activation function. The system generates a supervisory signal by monitoring the output error, solves the Fisher information matrix and KL scatter online via the analog computing unit, and finally feeds back to the α modulator and weight update circuit to form a closed-loop learning regulation.

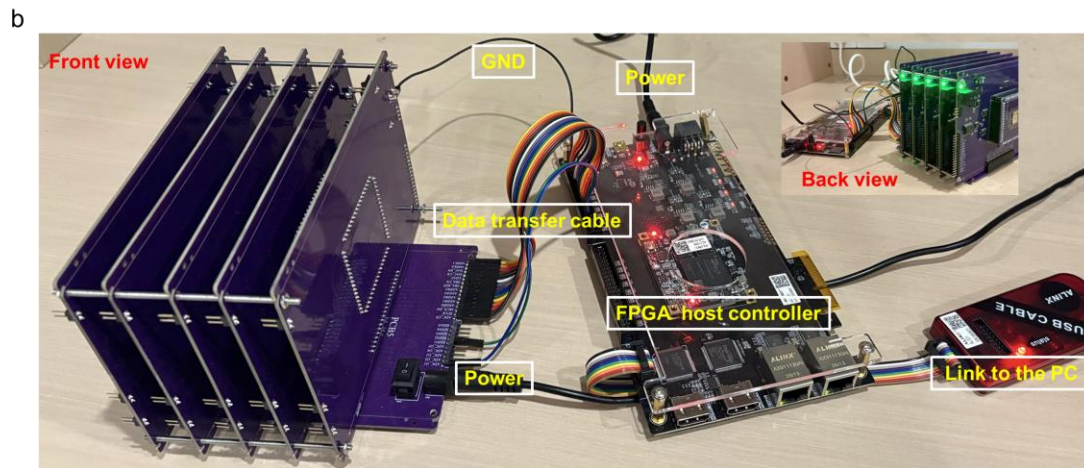
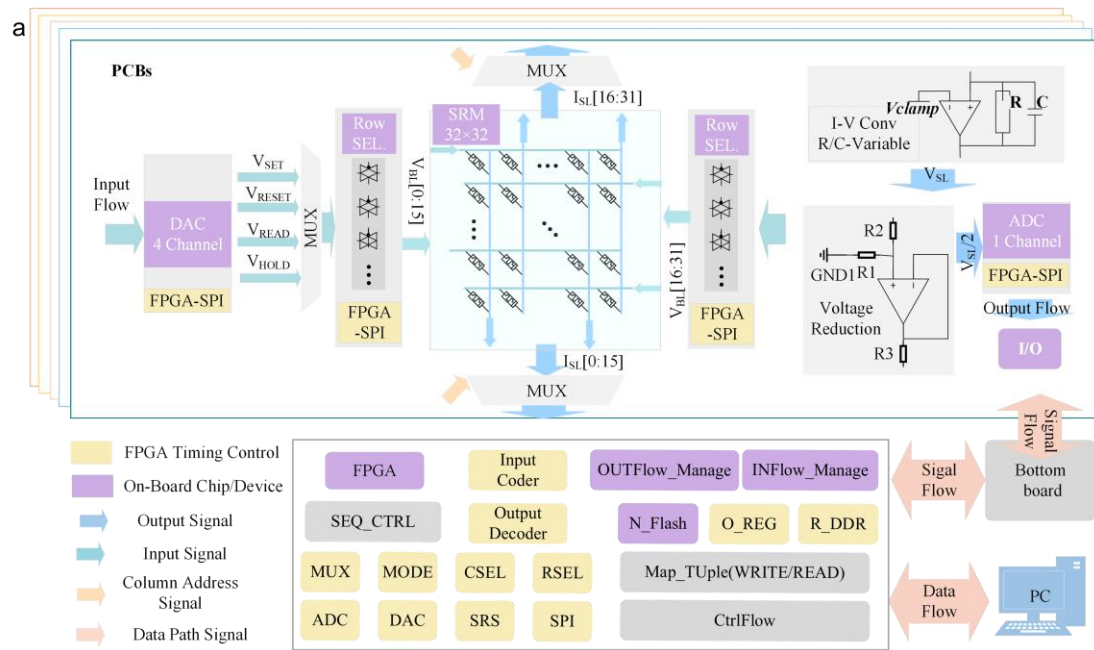


Fig.S21 Functional block diagram of the EHLL system. (A) The FPGA host controller supports data and information interaction between the PC and SPCs units, and generates corresponding control signals for all operations. The N_Flash numerical memory stores content from the personal computer (PC), including neural network weights, input numerical data, and calculation result data. The O_REG (operation register) holds individual data of mapping tuples. The R_DDR (result double data rate buffer) stores decoded convolution results. The underlying drivers include driver modules for ADC, DAC, SRS (shift register), and SPI (serial peripheral interface), as well as modules that control all MUXs to implement CSEL (column selection), RSEL (row selection), and mode selection. The bottom board with voltage and digital signal conversion circuits is used for the connection of other boards. (B) Physical prototype diagram of the EHLL system.

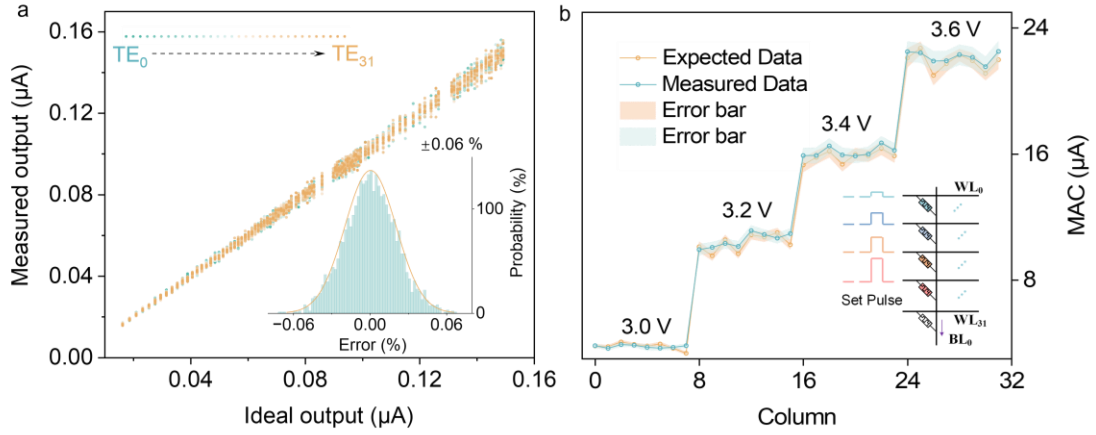


Fig.S22 Results of MAC operation performed by the SSM array. (A) The statistics of automatic quantized output and ideal output of 32 units during 200 consecutive set pulses were measured continuously. The ideal output value is obtained by applying a bias voltage to the electrode at the array intersection of the device unit. Each point represents one device. All read pulses adopt 1.0 V and 10μs (the inset shows the quantization error statistics). (B) The array is divided into unit regions with 8 columns as a unit. Four amplitudes of set pulses are used to program the unit regions respectively, with fixed pulse width (50 μs) and read pulse (1 V, 10 μs). The difference between the MAC operation result and the analog value reflects the training accuracy and feasibility.

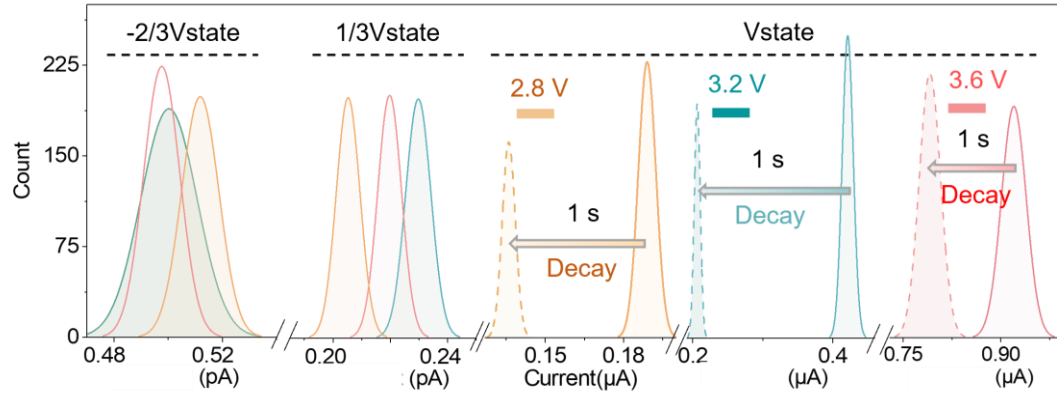


Fig.S23 Effect of voltage amplitude on weight relaxation of 1Kb array. Probability distributions of weights using the $1/3V$ programming scheme for 1024 devices in one array, under 2.8 V, 3.2 V, and 3.6 V programming pulses, and their corresponding values after a 1 s decay period. Experimental data follow a Gaussian distribution (dashed lines).

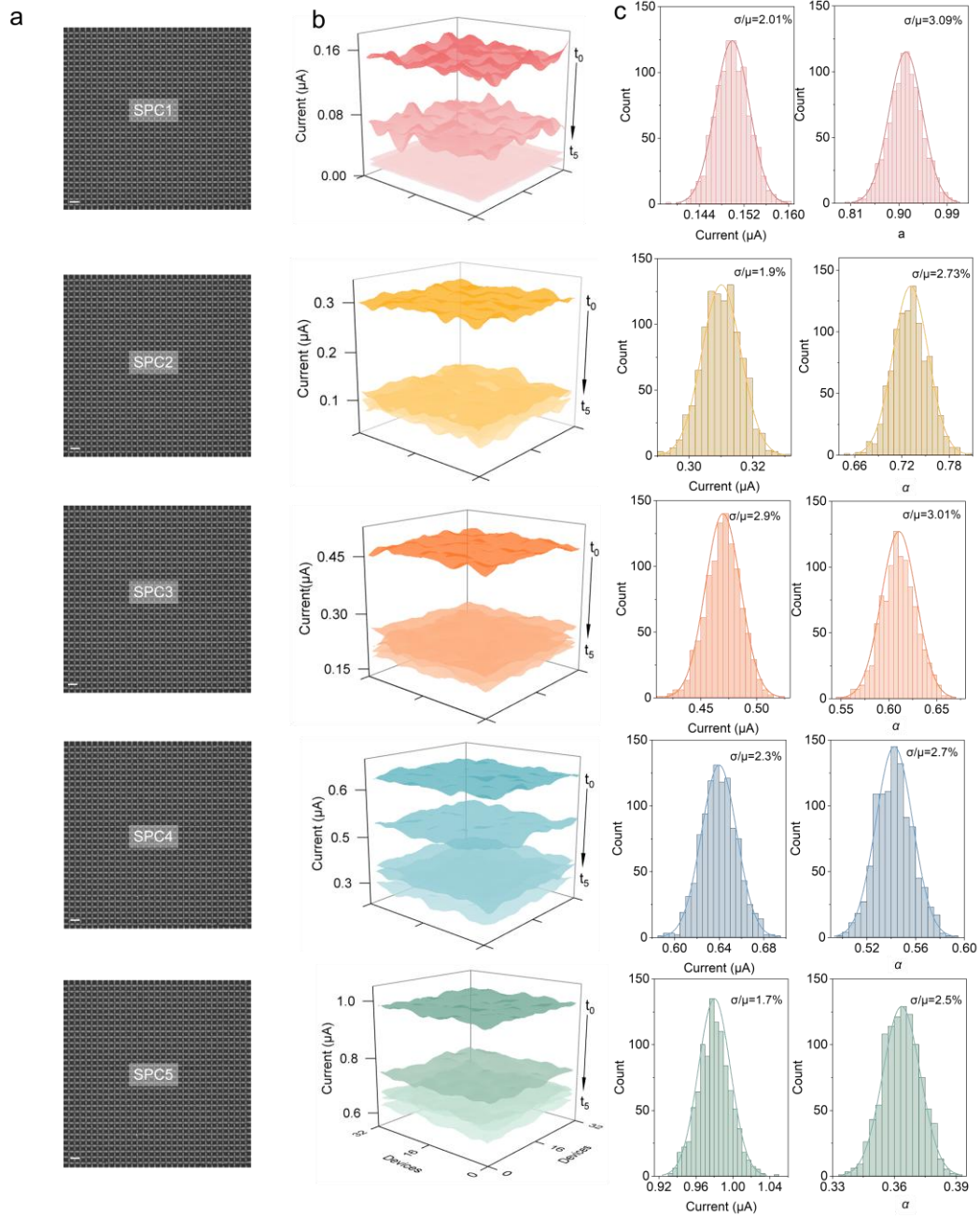


Fig.S24 Active forgetting characteristics of 5 SPCs. (A) The SEM image correspond to five 1Kb arrays with distinct forgetting rates, along with the array-level weight decay processes and the corresponding statistical data of weights and forgetting rates (Scale bar: 2 μ m). (B) Five 1Kb arrays were individually programmed using 200 consecutive pulses, which featured varying pulse amplitudes (2.8 V for SPC1, 3.0 V for SPC2, 3.2 V for SPC3, 3.4 V for SPC4, and 3.6 V for SPC5) and a fixed pulse width of 50 μ s. The weight immediately after pulse programming is defined as the weight at t_0 . With the interval of 200 ms, the weight variations of the array during the decay process within 1 s were recorded. Subsequently, (C) the weight distribution at t_0 and the forgetting rate at 1 s (t_5) were statistically analyzed. We statistically

quantified the α values following t_5 . $\alpha=(wt_5-wt_0)/wt_0$.

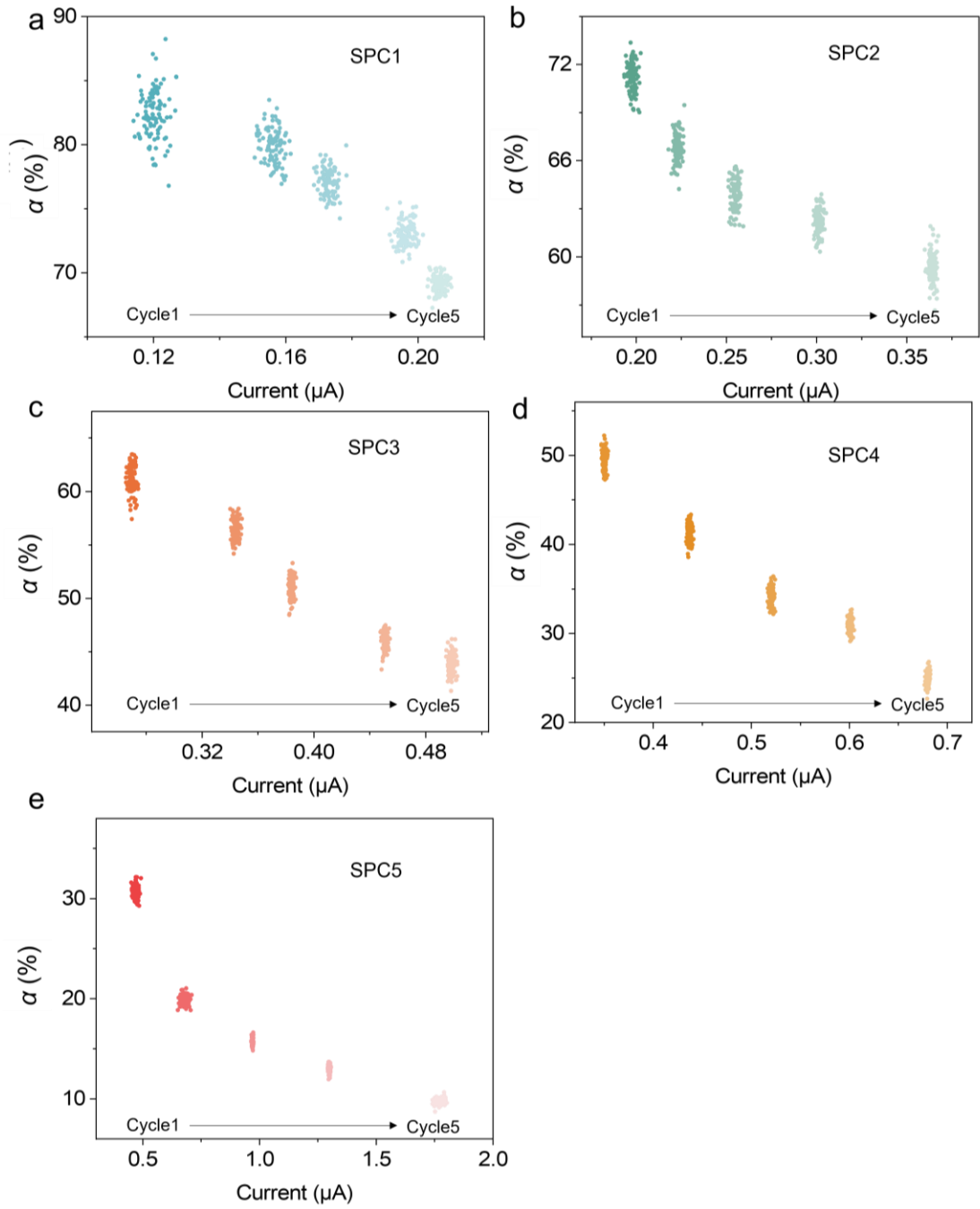


Fig.S25 Statistical analysis of the corresponding relationship between weight states and decay factor. 100 cyclic tests for 5 SPC cores over 5 programming cycles (corresponding to Fig. 3D); statistical analysis of potentiated weight and decay factor distributions in each cycle.

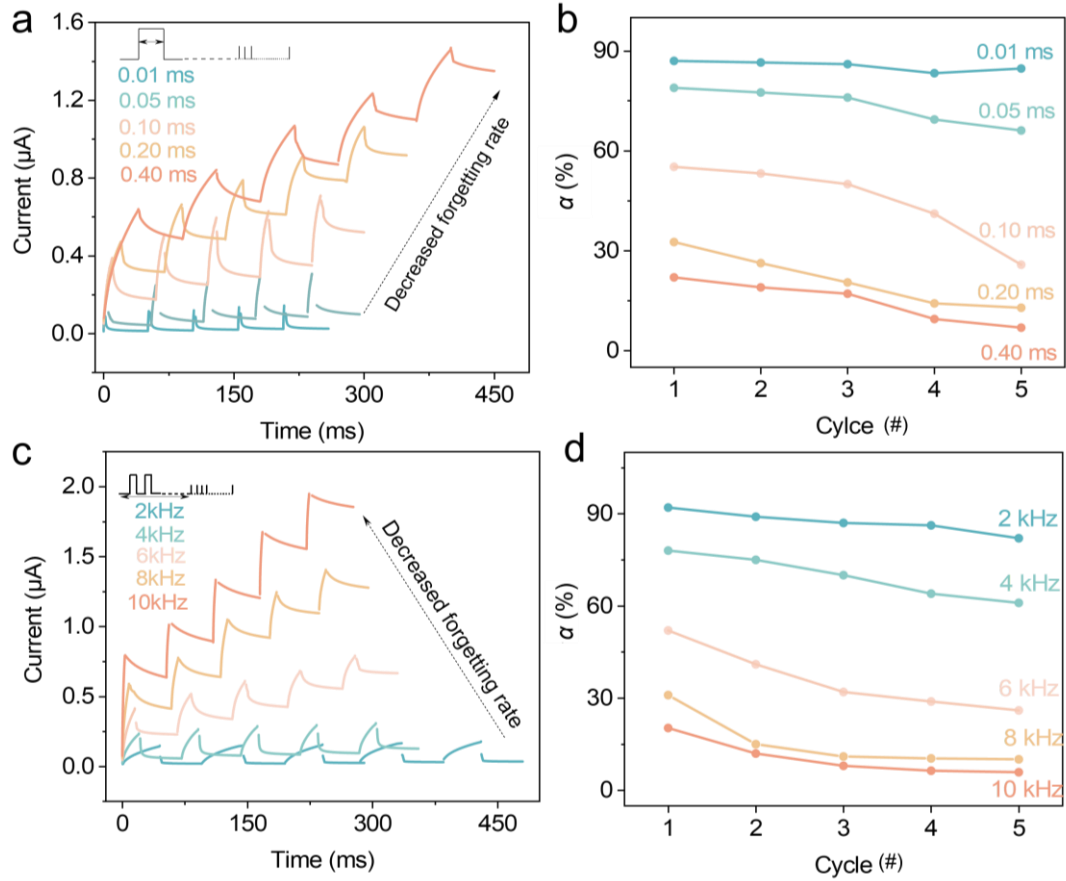


Fig.S26 Demonstration of the balance between LP and MS under different pulse widths and frequencies. **(A)** Weight programming for 5 cycles was performed on the 5 SSM arrays using 100 positive pulses and a 50 ms decay period, respectively. The pulse conditions used for testing were as follows: the positive pulses had an amplitude of 3 V, with widths of 0.01 ms (for SPC1), 0.05 ms (for SPC 2), 0.1 ms (for SPC 3), 0.2 ms (for SPC 4), and 0.4 ms (for SPC 5); the read pulses were 1 V with a width of 10 μs . **(B)** This shows the weight forgetting rate of the 5 SSM arrays over 5 consecutive cycles. A loop was run 100 times for the 5 cycles of the 5 SSM arrays, and the distributions of weights and weight forgetting rates across the five cycles were statistically analyzed. **(C)** Weight programming for 5 cycles was performed on the 5-layer SSM array using 100 positive pulses and a 50 ms decay period, respectively. The pulse conditions used for testing were as follows: the positive pulses had an amplitude of 3 V, a width of 50 μs , and frequencies of 2 Hz (for SPC 1), 4 Hz (for SPC 2), 6 Hz (for SPC 3), 8 Hz (for SPC 4), and 10 Hz (for SPC 5); the read pulses were 1 V with a width of 10 μs . **(D)** This shows the weight forgetting rate of the 5 SSM arrays over 5 consecutive cycles. A loop was run 100 times for the 5 cycles of the 5-layer SSM array, and the distributions of weights and weight forgetting rates across the five cycles were statistically analyzed.

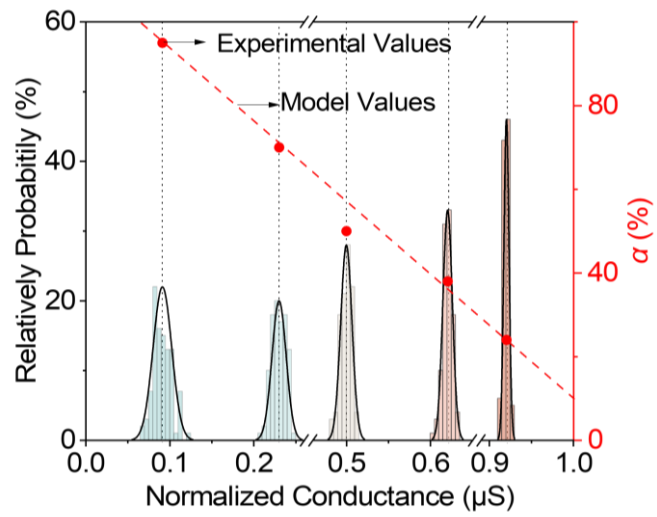


Fig.S27 Consistency between the modulation of the custom decay factor α and the weight update of SPCs.

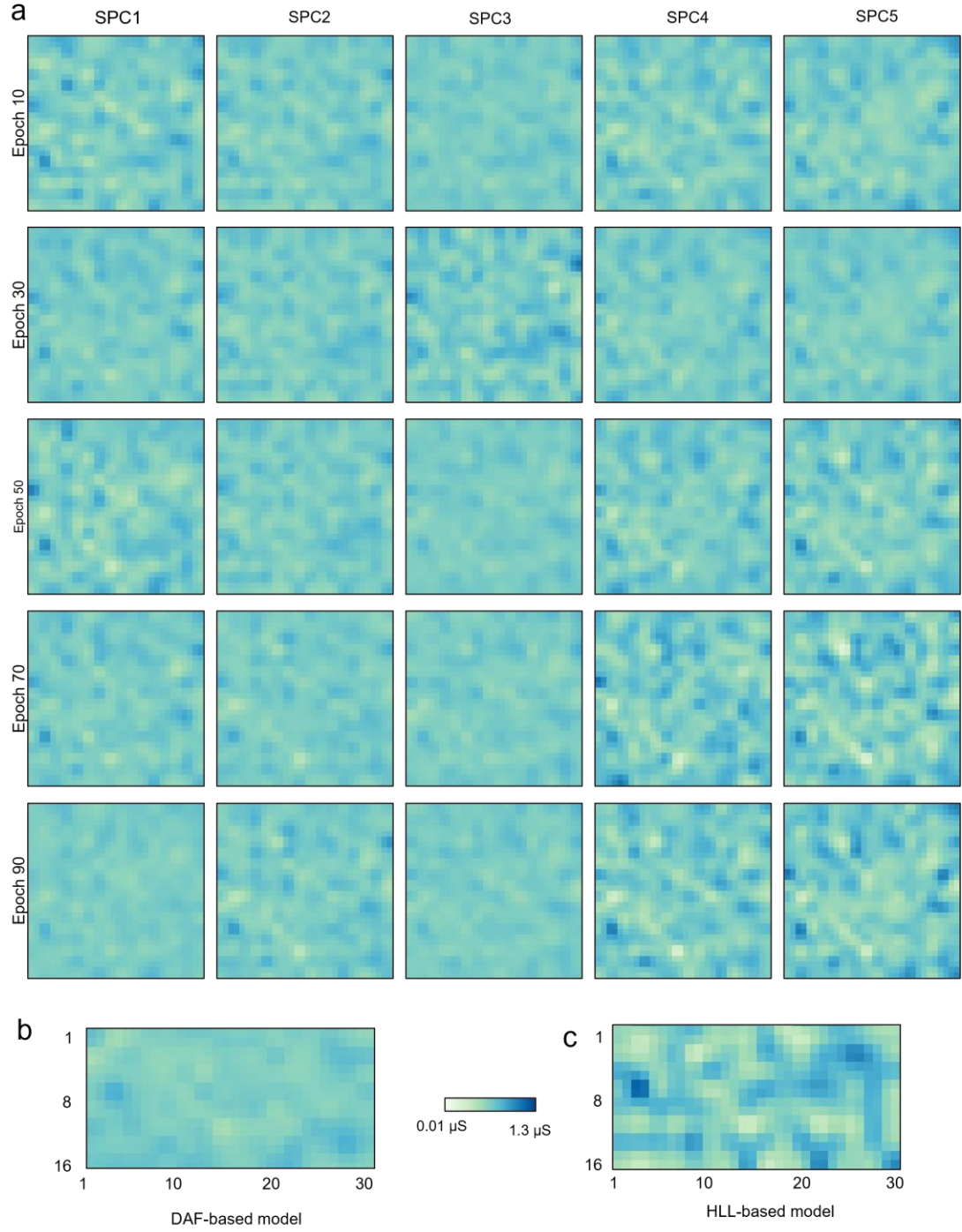


Fig.S28 Weight mapping after hybrid training. (A) The weight maps obtained from the HLL model learning S-CIFAR-100 dataset, which are mapped to the five SSM processing cores. Weight maps obtained via hybrid training on the first processing core of hardware-implemented (B) DAF-based and (C) HLL-based models across twenty sequential S-CIFAR-100 tasks.

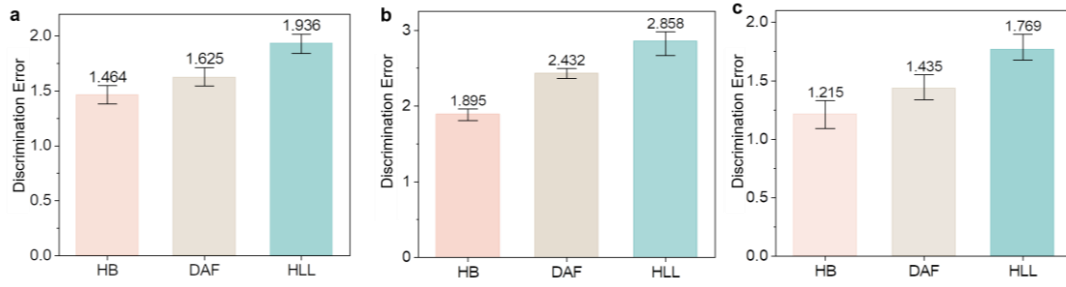


Fig.S29 Discrimination errors of three lifelong learning models based on SSM processing cores on three datasets. (A)S-CIFAR-100, (B)R-CIFAR-100, and (C) R-CIFAR-10/100. Discrimination error is calculated to measure the model's ability to distinguish between different tasks or domain data distributions[2]. Specifically, the error evaluates whether features of the input data belong to a particular task or domain by training a simple discriminator. The discriminator usually consists of a fully connected layer trained with a binary cross-entropy loss function, and its goal is to correctly categorize the features of the input data into the corresponding task or domain. In the testing phase, the discrimination error is calculated based on the average classification error rate of the discriminator on the test set, which reflects the difficulty of the model in distinguishing between different task or domain data distributions. The larger the error value, the smaller the difference in the feature distribution of different tasks or domains, indicating that the model is better able to map data with different distributions into a similar feature space, thus reducing the differences between domains.

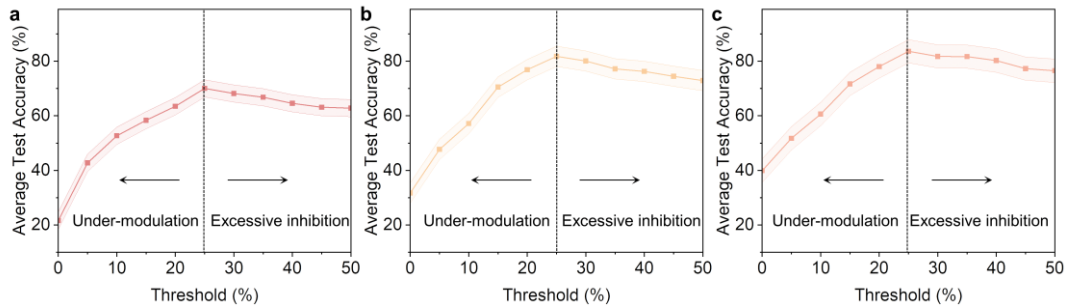


Fig.S30 Average test accuracy of the HLL model learned on three datasets. (A) S-CIFAR-100, (B)R-CIFAR-100, and (C) R-CIFAR-10/100, corresponding to different inhibitory thresholds for neuronal-aspect neuromodulation.

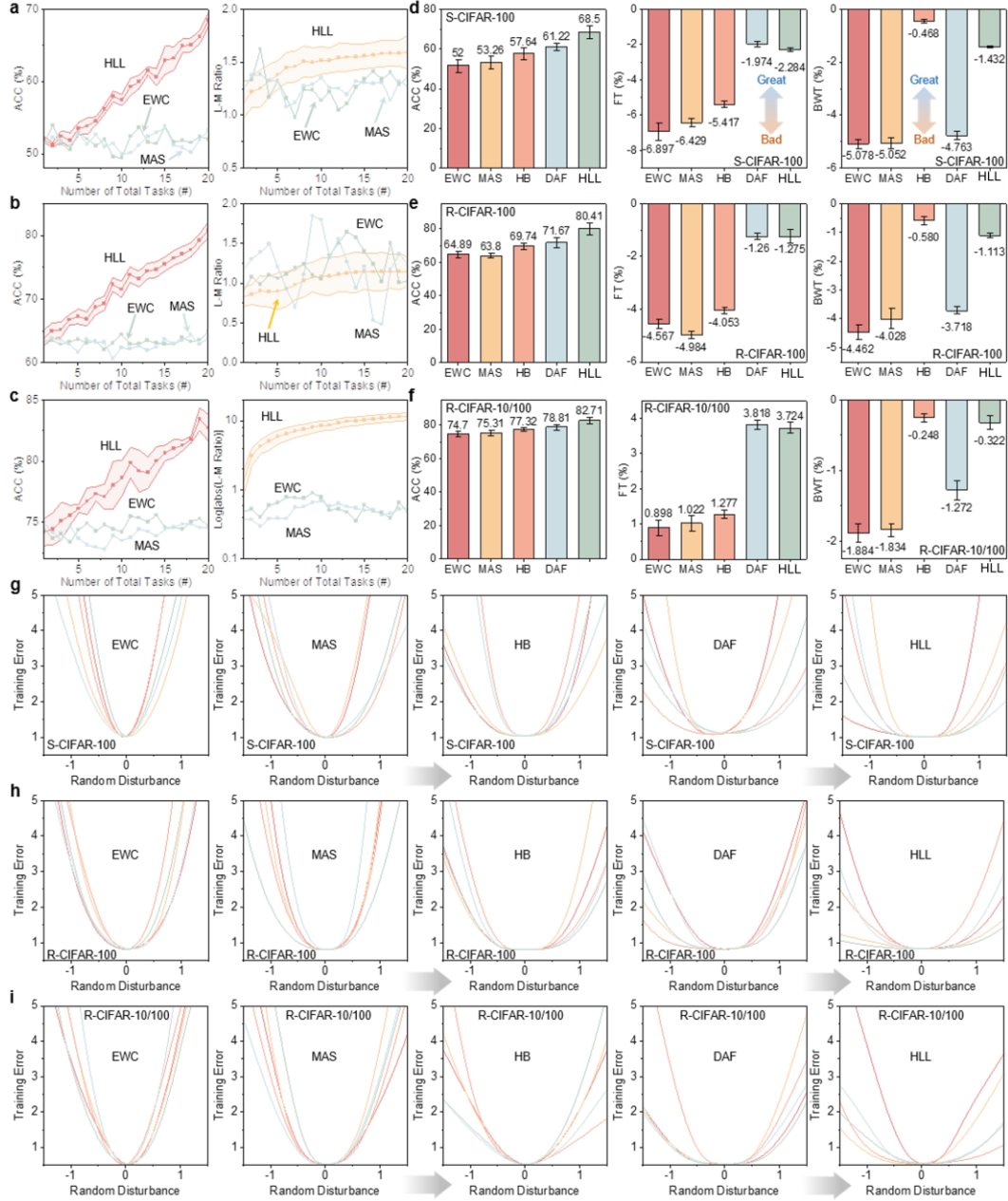


Fig.S31| Model generalization and robustness. General performance of the proposed SRM-based HLL model compared to conventional lifelong learning configuration. The different numbers of tasks corresponding to the average test accuracy (left) and the LP-MS ratio (right) obtained by training HLL, conventional E, and conventional MAS models on **(A)** S-CIFAR-100, **(B)** R-CIFAR-100, and **(C)** R-CIFAR10/100 datasets, respectively. Experimental results show that the average test accuracy of the proposed bionic HLL model increases steadily with the number of learned tasks, whereas that of the conventional regularization-based EWC and MAS models fluctuates sharply and decreases slightly. The core reason is that traditional models only maintain MS via parameter constraints, failing to address the stability-plasticity trade-off. The LP-MS ratio of the HLL model rises steadily and stabilizes with task input, driven by its inherent diversity mechanism toward optimal performance. The average test accuracy, LP, and

MS, as three key metrics of lifelong learning, obtained by training HLL, HB, DAF, conventional EWC, and conventional MAS models on **(D)** S-CIFAR-100, **(E)** R-CIFAR-100, and **(F)** R-CIFAR-10/100 datasets, respectively. On the S-CIFAR-100 dataset, standalone HB and DAF models outperform traditional ones significantly, while the HLL model achieves a more substantial accuracy gain by integrating DAF's high LP and HB's high MS for targeted metric optimization. Notably, some models show negative LP (old-task inhibition on new tasks) and MS (new-task-induced forgetting), whereas HLL attains lower forgetting via HB and stronger new-task adaptability via DAF. Similar trends appear on the R-CIFAR-100 dataset, and all models yield positive LP on the R-CIFAR-10/100 dataset—indicating shared features enable beneficial knowledge transfer, which HLL further enhances by mitigating knowledge conflicts through its AF property. The influence of the random disturbance obtained after training HLL, HB, DAF, conventional EWC and conventional MAS models on the **(G)** S-CIFAR -100, **(H)**. R-CIFAR-100 and **(I)** R-CIFAR-10/100 datasets over the training error in five parameter directions. Larger curvature (i.e., larger second-order derivatives) in certain directions indicates that the loss function varies drastically in that direction, and thus even applying smaller disturbance can lead to a significant increase in the training error, while directions with smaller curvature are relatively insensitive to disturbance, and the change in the training error is gentler. The error curves of all models exhibit distinct curvatures across different parameter dimensions, reflecting significant differences in their perturbation sensitivity: conventional EWC and MAS models suffer sharp rises in training error even under minor perturbations, indicating poor stability. In contrast, the bionic HLL model amplifies such parameter sensitivity differences via the diversity of multiple SSM-based learners and the stability of dual-sided neuromodulation, thus achieving a global balance between anti-perturbation robustness and task adaptability. These performance discrepancies stem from the fundamental differences in loss landscape exploration and stability-plasticity trade-off handling: traditional models rely on regularization to constrain parameter update paths, resulting in larger loss landscape curvatures and susceptibility to perturbations; the HLL model dynamically adjusts the parameter space exploration range, reduces the overall loss landscape curvature, and disperses perturbation impacts through multi-core prediction diversity, ultimately yielding a smoother error variation trend.

Captions for Video 1

The movie demonstrates that the EHLL system achieves continual learning across 20 tasks via the hardware- native adaptability of plasticity-stability balance. In contrast, a single biological model (DAF/HB) exhibits significant catastrophic forgetting when performing continual learning over 20 tasks.

Reference

- 1 Zhang, G. et al. Self-rectifying memristors with high rectification ratio for attack-resilient autonomous driving systems. *Nat. Commun.* **16**, 5759 (2025).

- 2 Zhang, T. et al. A brain-inspired algorithm that mitigates catastrophic forgetting of artificial and spiking neural networks with low computational cost. *Sci. Adv.* **9**, 2947 (2023).
- 3 Wang, L. et al. Incorporating neuro-inspired adaptability for continual learning in artificial intelligence. *Nat. Mach. Intell.* **5**, 1356-1368 (2023).
- 4 Rissanen, J.J. Fisher information and stochastic complexity. *IEEE Trans. Inf. Theory.* **42**, 40-47 (1996).
- 5 Shine, J.M. et al. Computational models link cellular mechanisms of neuromodulation to large-scale neural dynamics. *Nat. Neurosci.* **24**, 765-776(2021).
- 6 Yao, P. et al. Fully hardware-implemented memristor convolutional neural network. *Nature.* **577**, 641-646 (2020)
- 7 Lomonaco V, FilliatD, et al. Don't forget, there is more than forgetting: new metrics for Continual Learning, arXiv preprint arXiv:1810.13166, 2018. [Online]. Available: <https://arxiv.org/abs/1810.13166>
- 8 Kemker R, McClure M, et al. Measuring Catastrophic Forgetting in Neural Networks, Proceedings of the AAAI Conference on Artificial Intelligence, vol. 32, no. 1, 2018, doi: 10.1609/aaai.v32i1.11651.
- 9 Tuytelaars T, Tolias A S, Three types of incremental learning, Nature Machine Intelligence, vol. 4, no. 12, pp. 1185-1197, 2022, doi: 10.1038/s42256-022-00568-3.
- 10 Feng, Y. et al. Memristor-based storage system with convolutional autoencoder-based image compression network. *Nat. Commun.* **15**, 1132 (2024).