

Intrinsic Preservation of Plasticity in Continual Quantum Learning

Yu-Qin Chen^{1,*} and Shi-Xin Zhang^{2,†}

¹*Graduate School of China Academy of Engineering Physics, Beijing 100193, China*

²*Institute of Physics, Chinese Academy of Sciences, Beijing 100190, China*

 (Received 19 January 2026; revised 24 March 2026; accepted 21 May 2026; published 6 July 2026)

Artificial intelligence in dynamic, real-world environments requires the capacity for continual learning. However, standard deep learning suffers from a fundamental issue: loss of plasticity, in which networks gradually lose their ability to learn from new data. Here we show that quantum learning models naturally overcome this limitation, preserving plasticity over long timescales. We demonstrate this advantage systematically across a broad spectrum of tasks from multiple learning paradigms, including supervised learning and reinforcement learning, and diverse data modalities, from classical high-dimensional images to quantum-native datasets. Although classical models exhibit performance degradation correlated with unbounded weight and gradient growth, quantum neural networks maintain consistent learning capabilities regardless of the data or task. We identify the origin of the advantage as the intrinsic physical constraints of quantum models. Unlike classical networks where unbounded weight growth leads to landscape ruggedness or saturation, the unitary constraints confine the optimization to a compact manifold. Our results suggest that the utility of quantum computing in machine learning extends beyond potential speedups, offering a robust pathway for building adaptive artificial intelligence and lifelong learners.

DOI: [10.1103/ry52-85ll](https://doi.org/10.1103/ry52-85ll)

I. INTRODUCTION

Recent advances in artificial intelligence (AI) [1,2], particularly in large language models, rely on scaling static architectures on massive fixed datasets [3–5]. However, the deployment of these systems in the real world reveals a critical limitation: the world is non-stationary, and data distributions evolve continuously over time. It is computationally prohibitive to retrain large-scale models from scratch to accommodate every distributional shift and new knowledge [6]. Consequently, the paradigm of continual learning (CL) emerges [7–9], where models learn sequentially from a stream of tasks. The central challenge of CL lies in the trade-off between *stability* and *plasticity*. Stability requires the model to keep proficiency on previously learned tasks, preventing catastrophic forgetting [10–12], whereas plasticity demands the flexibility to efficiently learn new knowledge from fresh data streams. While the community has dedicated significant effort to the former challenge of stability, a more subtle problem has recently been identified regarding the latter, dubbed

loss of plasticity [13]. Recent studies demonstrate that standard deep learning methods not only forget; they also progressively lose their learnability, with the performance on subsequent tasks decaying continuously [13–20].

Recently, quantum machine learning (QML) has emerged as a promising alternative learning paradigm, leveraging quantum principles to process information in high-dimensional Hilbert spaces [21]. Theoretical foundations have established the expressive power and generalization capabilities of quantum models [22–26]. Specifically, variational quantum circuits (VQCs) and quantum neural networks (QNNs) have been successfully applied to diverse tasks, including classification [27,28], generative modeling [29–31], and quantum-native data processing [32,33]. Furthermore, recent studies explore sophisticated architectures, optimization strategies, and utilize the intrinsic information properties to enhance QML performance and resilience [34–45]. In the context of continual learning, initial investigations focused primarily on the stability aspect of CL. Previous studies demonstrated that QML approaches based on variational quantum circuits (VQCs) cannot naturally mitigate catastrophic forgetting, exhibiting forgetting dynamics similar to classical counterparts under standard training settings [46,47]. However, the plasticity nature of quantum models remains unexplored. Therefore, it is an urgent and relevant question whether the unique structure of quantum neural networks can offer a decisive plasticity advantage in maintaining learnability over time. A positive answer would suggest that quantum computing

* Contact author: yqchen@gsaep.ac.cn

† Contact author: shixinzhang@iphy.ac.cn

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

offers an interesting mechanism toward adaptive intelligence, enabling systems that can adapt to dynamic environments without the huge cost of retraining or the risk of classical plasticity loss.

In this work, we address the critical open question through a systematic comparison of classical multi-layer perceptrons (MLPs) and deep quantum neural networks (QNNs), which are the representative models for classical and quantum AI, respectively. We evaluate their long-term adaptability across a variety of continual learning paradigms including supervised classification and reinforcement learning. The evaluation uses both industrial standard classical benchmarks and quantum-native datasets. Unlike the small-scale proof-of-concept demonstrations typical of prior QML research, our framework involves training very deep parameterized quantum circuits end-to-end across thousands of sequential tasks. This represents a computational workload orders of magnitude more challenging than standard benchmarks, validating QML plasticity in a more realistic and scalable regime. Our analysis reveals a qualitative difference in their learning behaviors. Classical models show loss of plasticity as expected with unbounded weight growth. On the contrary, as long as the quantum model is trainable, namely, capable of bypassing static trainability hurdles such as barren plateaus, its unitary character naturally preserves the model’s plasticity, regardless of the training duration. We attribute this plasticity to the unitary nature of evolution that confines optimization to a compact parameter manifold. While the compactness of quantum parameter manifolds is a well-known physical property, establishing its nontrivial algorithmic consequence for mitigating plasticity loss constitutes a central conceptual contribution of this work. By preventing the saturation dynamics that paralyzes classical networks, our results establish a distinct structural advantage for QNNs. This shifts the focus from computational speed to inherent adaptability, identifying quantum machine learning as a promising ingredient for realizing truly adaptive, lifelong learning agents.

II. PRESERVED PLASTICITY IN SUPERVISED LEARNING

To systematically evaluate the plasticity of QNNs versus classical counterparts, we first deploy them on the standard continual learning benchmark of Permuted MNIST. We note that more advanced architectures, such as convolutional neural networks and ResNet [48], have already been shown to suffer from severe plasticity loss in this setting [13]. In our experimental setup, we generate a sequence of 1000 tasks, where each task requires 10-fold classification on handwritten MNIST digits subject to a task-specific random pixel permutation [Fig. 1(a)]. We compare a classical MLP with ReLU activation against a VQC constructed using $SU(4)$ entangling gates [49,50]. The MLP is configured with hidden widths (NW) ranging from

4 to 256 neurons, while the QNN utilizes varying circuit depths (DP) of 2 to 16 layers. For the QNN output, we employ a probability-based readout strategy, i.e., the logarithmic measurement probabilities of the first ten bitstrings are mapped to class logits for softmax function output, analogous to the classical output layer. Both models are trained using gradient-based optimization on each task’s training set and evaluated on a held-out test set.

As shown in Fig. 1(b), classical MLPs exhibit a continuous degradation in performance. Despite the sufficient learning capacity for each individual task, the test accuracy decays significantly as the task sequence progresses. On the contrary, QNNs maintain a stable accuracy throughout the entire CL sequence. Figure 1(c) quantifies this difference: while MLPs suffer a large accuracy drop, QNNs exhibit negligible degradation. We verify that these divergent learning patterns persist across different model sizes and training configurations.

To probe the underlying mechanism behind this divergence, we monitor the internal dynamics of the networks during CL. We compute the relative L_2 norms of the weight matrices ($\|\theta\|_2$) and the gradients ($\|\nabla_{\theta}\mathcal{L}\|_2$) at each task [Figs. 1(d) and 1(e)]. In the classical case, we observe a monotonic, unbounded growth in weight magnitude. This weight explosion severely limits plasticity by pushing neurons into inactive regimes (e.g., dead ReLUs), leading to a corresponding collapse in learning capacity and effective rank. In direct contrast, the parameters of our $SU(4)$ -based QNN are rotation angles on a compact Lie group manifold. This unitary constraint ensures that the parameter space remains bounded and periodic. Consequently, both the weight and gradient norms of the QNN are stable, preventing the saturation dynamics that paralyze classical ones.

To confirm that plasticity preservation is not unique to the $SU(4)$ quantum *Ansatz*, we also test a restricted hardware-efficient *Ansatz* [51] for QNN. We find identical plasticity preservation (see Appendix C), confirming that the advantage is a property of the quantum model’s nature rather than a specific circuit architecture.

We extend our evaluation to a more challenging visual domain using the Split CIFAR-100 benchmark [Fig. 2(a)]. We build a sequence of 3000 binary classification tasks, where each task is to distinguish between two randomly sampled classes, such as apple versus orange. Inputs are encoded into quantum states via amplitude encoding after L2-normalization, and the output logit of the QNN is determined by the expectation value of a Pauli Z operator on the last qubit. Notably, we utilize deep QNNs with up to 30 layers (more than 4000 quantum trainable parameters) on this real-world dataset, demonstrating the scalability of the finding beyond typical toy-model experiments.

Consistent with the MNIST results, standard MLPs suffer a catastrophic loss of plasticity, with accuracy collapsing toward random guessing after approximately

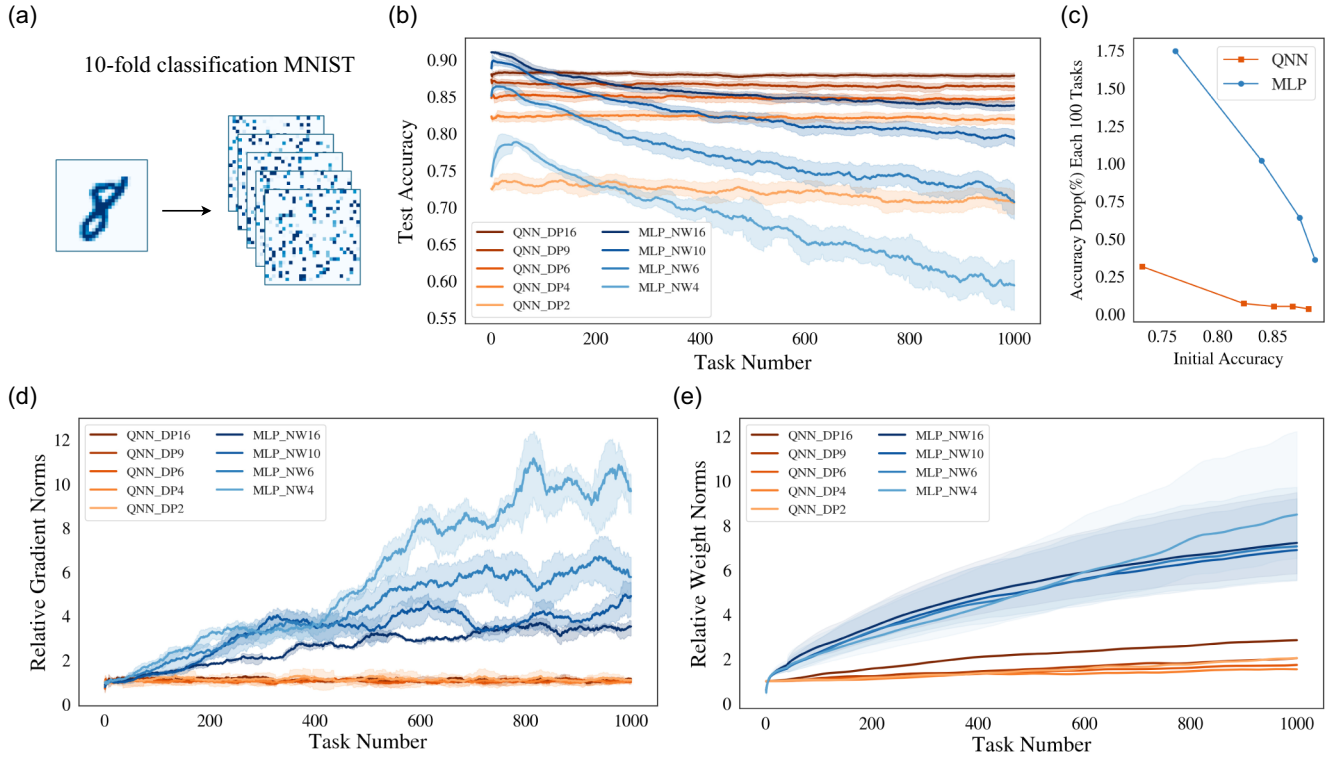


FIG. 1. Quantum neural networks maintain plasticity while classical neural networks fail in deep continual learning on permuted MNIST. (a) Illustration of the permuted MNIST task for continual learning. A sequence of 10-fold classification tasks is generated by applying different random permutations to the pixels of the MNIST digits, forcing the model to continuously adapt to new input distributions. (b) Test accuracy as a function of the task number for QNNs of varying circuit depths (DP) and MLPs of varying network widths (NW). QNNs maintain a stable test accuracy across the whole task sequence, demonstrating robust plasticity. In contrast, all MLP variants exhibit a significant and continuous degradation in performance, indicating a severe loss of plasticity. (c) Average accuracy drop against the model's initial accuracy averaged over the first 200 tasks. MLPs show a much higher rate of plasticity loss, and it is more severe for models with worse initial performance. (d) Relative gradient norms and (e) relative weight norms during the continual learning process. The norms are normalized by the average norm across the first 10 tasks. For MLPs, both gradient and weight norms grow unboundedly over time, correlating with their performance decline. For QNNs, these norms remain stable and bounded, providing a mechanism explanation for the preserved plasticity. All curves show the moving average over 40-task windows, with shaded areas representing the standard deviation of the underlying data within each window.

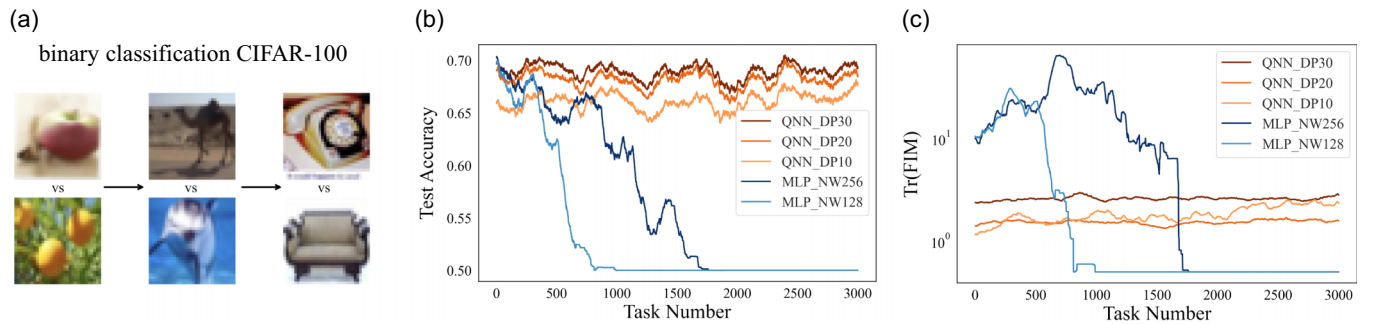


FIG. 2. Preserved plasticity of QNNs in continual learning of real-world image classification tasks. (a) Illustration of the sequential binary classification task constructed from the CIFAR-100 dataset. At each task, two classes are randomly sampled from the 100 available classes to form a new binary classification problem. This setting challenges the model's ability to continuously adapt to new visual concepts. (b) Test accuracy as a function of task number for QNNs of varying circuit depths (DP) and MLPs of varying network widths (NW). Similar to the results on permuted MNIST, QNNs maintain a stable test accuracy throughout the 3000 tasks. In stark contrast, MLPs suffer a catastrophic loss of plasticity, with their performance rapidly decaying to guess level (50% accuracy). (c) Mechanism insight by the trace of the Fisher Information Matrix, a measure of the model's effective learning capacity. Tr(FIM) for QNNs remains stable, indicating that their ability to learn is preserved. For MLPs, the Tr(FIM) collapses, quantitatively demonstrating that the model has lost its capacity to effectively update its parameters, which explains the performance decay in panel (b). All curves show the moving average over 150-task windows.

1000 tasks Fig. 2(b)]. Crucially, for a fair comparison, both models are trained using exactly identical optimization hyperparameters and learning schedules, ensuring that the quantum advantage originates from the *Ansatz* geometry rather than training hyperparameters. To further isolate the structural benefit of quantum models, we also evaluate an L2-regularized classical MLP with weight decay set to 10^{-4} . Even with explicit regularization, the classical model still collapses to random guessing accuracy within the same time frame. In stark contrast, QNNs demonstrate superior robustness, maintaining high classification accuracy throughout the extensive 3000-task process.

We also verify the robustness of this advantage on hardware-realistic noisy intermediate-scale quantum (NISQ) devices by simulating depolarizing noise ($p = 5 \times 10^{-4}$). The noisy QNN maintains the learnability, confirming that the geometric preservation of plasticity is robust to small quantum errors.

We further train a classical MLP with a periodic sine activation function [52] on Split CIFAR-100. While this introduces bounded activations similar to quantum amplitudes, the weights themselves remain unbounded in Euclidean space. The MLP using the $\sin(x)$ activation function (Sin-MLP) from severe plasticity loss and overfitting, as the unbounded weight growth drives the optimization landscape into high-frequency oscillation regimes (see Appendix D). This confirms that true plasticity preservation requires the compact parameter manifold inherent to the quantum *Ansatz*, not merely bounded output functions.

To provide a rigorous theoretical metric for learnability, we compute the trace of the Fisher information matrix (FIM) on a fixed probe dataset throughout the training [Fig. 2(c)] [53]. While the FIM trace is an information-theoretic measure of a model’s capacity, we specifically employ the empirical FIM in the main text, which represents the curvature of the loss surface with respect to the labels of the current batch. This metric serves as a direct proxy for the model’s update agility. For MLPs, $\text{Tr}(\text{FIM})$ undergoes a rapid collapse, quantitatively demonstrating that the effective dimensionality of the optimization landscape shrinks over time. The classical network physically possesses parameters, but they become functionally dead. QNNs instead maintain a stable $\text{Tr}(\text{FIM})$, indicating that their effective learning capacity remains constant regardless of the time in the continual learning history. We provide a more fundamental analysis using the model FIM (γ from the model predicting instead of data ground truth in the empirical FIM case) and rank analysis in the Appendix H, confirming that both metrics capture the same structural divergence in plasticity.

We further elevate these empirical findings to a rigorous theoretical analysis in Appendix B. By analyzing the asymptotic geometry of the optimization landscape, we prove a fundamental difference in learnability between classical and quantum models. For classical networks trained under cross entropy, we derive that as weight

magnitudes diverge in Euclidean space ($\lambda \rightarrow \infty$), the trace of the FIM collapses exponentially: $\lim_{\lambda \rightarrow \infty} \text{Tr}(\text{FIM}) \rightarrow 0$. In sharp contrast, for QNNs, we utilize the properties of Haar integration to prove that the expected learning capacity is strictly confined within a stable regime, bounded away from zero (preventing saturation) and also bounded from above (preventing extreme learning sensitivity as chaos). Consequently, the corresponding theorem implies that the quantum learning capacity is invariant to weight inflation, ensuring the trace of the FIM never decays due to the growing weights.

III. ROBUSTNESS IN DEEP REINFORCEMENT LEARNING

We next investigate whether the plasticity advantage extends to deep reinforcement learning (RL), where the data distribution shifts dynamically due to the agent’s

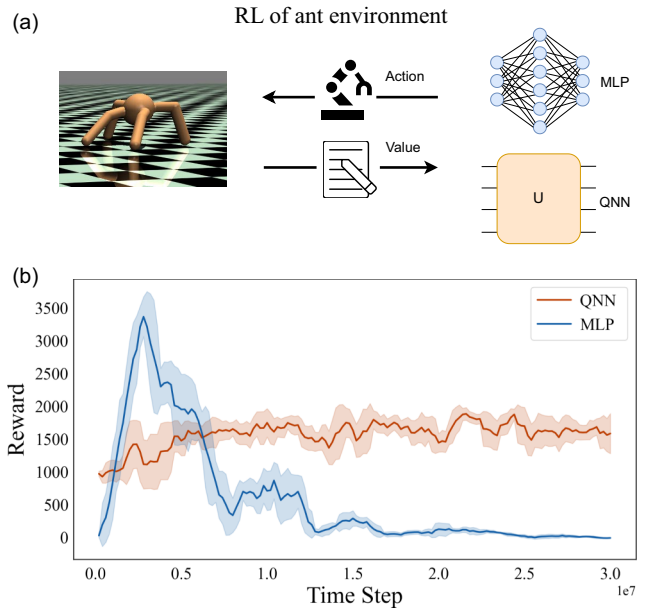


FIG. 3. Quantum-enhanced agents maintain learning plasticity in a standard reinforcement learning benchmark. (a) Schematic of the RL setup. An agent, using either a classical MLP or a quantum QNN as its function approximator for the policy and value networks, interacts with the Ant-v4 environment. The agent’s goal is to learn a motion policy that maximizes reward for efficiently moving forward. (b) Evaluation reward (evaluated at each 200,000 steps) as a function of time step during training with the PPO algorithm in a stationary environment. The agent with an MLP-based policy initially learns a strong policy, reaching a high reward. However, it subsequently suffers a catastrophic collapse in performance, demonstrating a severe loss of plasticity even without an explicit change of the task. Unlike the classical agent, the QNN-based agent learns more steadily and maintains a stable reward policy over the long term, showcasing its inherent ability to preserve plasticity in the dynamic context of RL. Curves show the moving average over a 5-evaluation reward points window, with shaded areas representing the standard deviation of the underlying data within each window.

evolving policy. We employ the high-dimensional Ant-v4 environment [Fig. 3(a)]. In this continuous control task, the agent must coordinate four legs to maximize a composite reward signal, which incentivizes forward distance traveled and survival, while penalizing excessive control torques and contact forces.

We compare a standard PPO agent using MLP networks against a hybrid quantum-PPO agent where the feature extractor is replaced by a 9-qubit, 12-layer quantum circuit with $SU(4)$ parameterized gates. The measurement probabilities of the quantum circuit are transformed by a scaled Tanh function to generate a dense feature vector, which is then projected via learnable linear heads to produce the stochastic action distribution and the scalar value estimate. Both agents are trained for 30 million time steps across 16 parallel environments, demonstrating the remarkable scalability of our findings.

As shown in Fig. 3(b), the classical agent initially learns a high-performing policy but suffers a severe policy collapse in the later stages of training, a phenomenon identified as a symptom of plasticity loss in RL [17,18]. However, the quantum agent shows no such collapse. While its initial learning curve is more gradual, it maintains a stable, high-reward policy over much longer time steps. This result suggests that the intrinsic plasticity of quantum models acts as a natural regularizer against the instability of long-horizon RL, effectively stabilizing the policy optimization.

IV. DEMONSTRATION ON QUANTUM-NATIVE DATA

Finally, we analyze the performance on a quantum-native continual learning task to confirm the universality of the observed advantages. We generate a sequence of 2000 binary classification tasks based on the eigenstates of a one-dimensional XXZ spin chain with varying anisotropy parameter Δ [Fig. 4(a)]. For this task, the QNN naturally operates in the same Hilbert space as the data input, with predictions derived from the expectation value $\langle Z \rangle$ of the first qubit.

Figure 4(b) reveals a profound result: classical MLPs still lose plasticity on the quantum dataset, which extends our findings to different data modalities. In contrast, QNNs demonstrate a dual advantage unique to the quantum-native regime. First, it learns better: the deeper QNN variant achieves significantly higher accuracy than the MLP, benefiting from the inductive bias alignment between the quantum model and the quantum data structure. Second, it learns longer: unlike the MLP, which degrades over time, the QNN preserves this high performance with zero degradation over 2000 tasks. The synergy between superior representational power and indefinite plasticity positions QML as a unique powerful paradigm for processing varying quantum information streams.

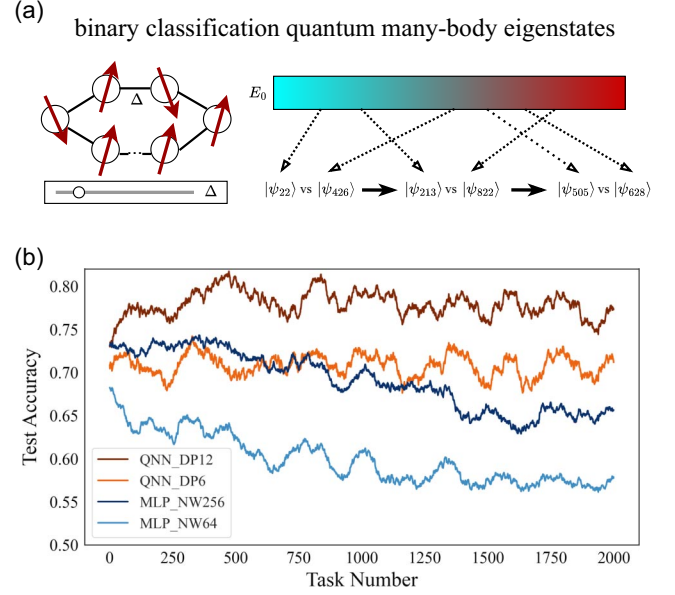


FIG. 4. QNNs exhibit superior plasticity on continual learning tasks with quantum native data. (a) Schematic of the quantum native continual learning task. The dataset consists of many-body eigenstates from the one-dimensional XXZ Hamiltonian with periodic boundary conditions, generated by varying the anisotropy parameter Δ . Each task in the continual learning sequence is a binary classification problem to distinguish between two sets of randomly sampled eigenstates from the full energy spectrum: the i th eigenstate $|\psi_i\rangle$ and the j th eigenstate $|\psi_j\rangle$. (b) Test accuracy as a function of task number for QNNs and MLPs of varying model sizes. QNNs, particularly the deeper variant (DP12), successfully learn and maintain a high classification accuracy across the sequence of 2000 tasks. In contrast, the MLPs show a clear degradation in performance, indicating that the loss of plasticity is a fundamental issue for classical models even for data with inherent quantum structure. This result demonstrates the robust plasticity of QNNs in their native domain. Curves show the moving average over a 100-task window.

V. DISCUSSION

Our findings suggest that the loss of plasticity is a pervading failure mode in deep learning due to the training in unbounded Euclidean spaces. In classical networks, the continuous accumulation of weight updates drives the model into saturation regimes where the effective rank of the representation collapses. Our counterfactual experiments with Sin-MLPs reveal that simply bounding the activation function is insufficient for classical learning models to maintain plasticity. The quantum *Ansatz*, by virtue of its unitary evolution and periodic parameterization on a compact manifold, naturally enforces the parameter space constraint. This nature ensures that the optimization landscape remains navigable indefinitely, as evidenced by the stable test accuracy and Fisher information spectrum during CL. Analytically, this contrast is detailed in Appendix B, where we define and rigorously

prove the asymptotic limits of the Fisher information spectrum for both architectures. It is important to emphasize that this dynamic advantage is conditional on initial trainability. Once in a trainable regime, the capacity does not dynamically decay.

While it is well established that quantum parameters reside on a compact periodic manifold, discovering a practical problem that fundamentally benefits from this geometric property has remained elusive. Our central innovation lies in bridging this gap. We identify the loss of plasticity in continual learning as the exact real-world scenario where this quantum constraint transforms into a decisive advantage. Thus, we demonstrate that quantum unitarity is not merely a physical rule but a structural solution to a fundamental pathology in artificial intelligence.

We remark that within the classical deep learning community, algorithmic interventions such as regenerative regularization [54], new parameter injection [55], and selective weight resetting (e.g., continual backpropagation) [13] have been proposed to mitigate plasticity loss. However, these approaches represent extrinsic patches: they rely on heuristic modifications to the training objective or invasive surgeries on the network state to counteract the model’s natural tendency toward rigidity. Such methods often introduce sensitive hyperparameters and can disrupt prior knowledge. The significance of our work lies in demonstrating that QNNs do not require these extrinsic interventions. Their plasticity is intrinsic, emerging directly from the QNN nature. This distinguishes the quantum approach as a first-principle structural solution rather than an engineering workaround, offering a more robust foundation for CL.

This insight suggests a redefinition of the strategic utility of QML. Historically, the field has been mainly focused on the search for computational speedups, typically quantified by asymptotic complexity. Our work instead identifies a distinct and perhaps more immediate axis of quantum advantage: *learning robustness*. While QNNs may not always train faster than classical counterparts, they learn longer and more reliably in non-stationary environments. This finding complements recent research demonstrating quantum advantages in *learning resilience* against data noise and unlearning efficiency [44], collectively suggesting that the geometric properties of quantum learning models offer a superior basis for adaptability. With learning robustness and resilience revealed, the evaluation on QML shifts the metric from mere *time-to-solution* to *quality-of-solution* over the model’s entire life cycle. This distinction is critical for real-time control systems and sustainable AI, where the train-reset-retrain cycle is ecologically prohibitive. Furthermore, on quantum-native tasks, we observe a dual advantage in terms of superior representational alignment and indefinite plasticity, which make QML the viable path for processing continuous streams of quantum information without degradation.

We conclude by noting the limitations and future directions of this research. Our study is conducted via exact simulations using TensorCircuit-NG [56]. The effect of quantum noise on continual quantum learning is worth further investigation. Besides, it is also interesting to explore whether we can design efficient quantum-inspired classical learning gadgets capable of maintaining plasticity like quantum models. Finally, we emphasize that while the quantum approach intrinsically preserves plasticity, it does not inherently solve the complementary challenge of stability. Both quantum and classical models in our study exhibit catastrophic forgetting of previous tasks when trained without explicit stability-preserving regularizers. A complete continual learning system would require augmenting the plastic quantum backbone with specialized stabilization techniques or memory mechanisms. This hybrid approach represents a promising frontier for realizing better AI that is both sufficiently stable to remember and plastic enough to evolve.

ACKNOWLEDGMENTS

This work was supported by the National Natural Science Foundation of China (No. 12504599 and No. 12574546), Quantum Science and Technology-National Science and Technology Major Project (No. 2024ZD0301700 and No. 2025ZD0300802), Science Challenge Project (No. TZ2025017), and the Chinese Academy of Sciences (No. XDB1680201 and No. YSBR-150).

DATA AVAILABILITY

The data that support the findings of this article are openly available [57].

APPENDIX A: METHODS AND PIPELINES

1. Experiment 1: Supervised continual learning on permuted MNIST

Task Formulation. We simulate a domain change learning scenario using the standard MNIST dataset. The experiment consists of a sequence of $T = 1000$ independent classification tasks. For each task t , a fixed random permutation matrix P_t is generated. The raw input images (28×28 pixels) are flattened into vectors $\mathbf{x} \in \mathbb{R}^{784}$. For task t , the input provided to the model is $\mathbf{x}'_t = P_t \mathbf{x}$. The 10-class target labels $y \in \{0, \dots, 9\}$ remain unchanged.

Data Preprocessing for Quantum Models. The inputs are zero-padded to map the 784-dimensional vector to a N -qubit Hilbert space ($d = 2^N$). We utilize $N = 10$ qubits ($d = 1024$). The 784-pixel vector is padded with zeros to form a vector of size 1024. This vector is then L_2 -normalized to unit magnitude to serve as a valid quantum state vector for amplitude encoding.

Model Architectures.

- **Classical MLP:** The network consists of an input layer ($d = 784$), a single hidden dense layer with

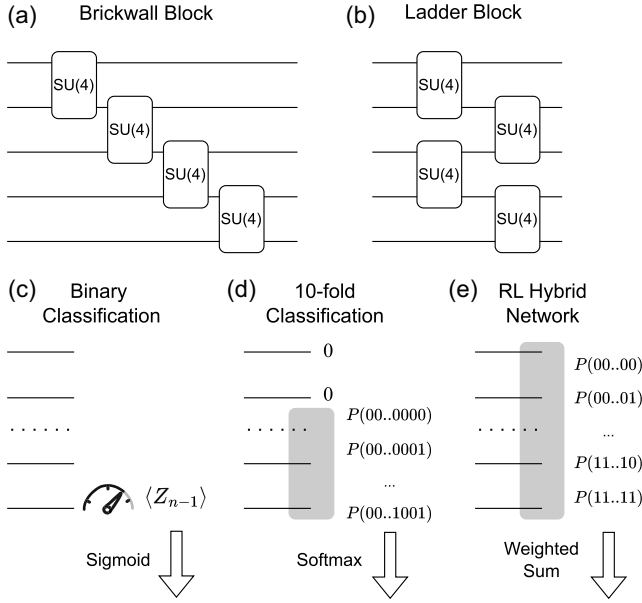


FIG. 5. Quantum neural network architectures and readout strategies employed in different experimental settings. (a), (b) Schematics of the variational *Ansatz* blocks used to construct the deep QNNs. Each block is composed of parametrized 2-qubit general unitary gates, $SU(4)$, acting on nearest neighbor qubits. (a) The Brickwall Block architecture applies gates to alternating pairs of qubits. (b) The Ladder Block architecture applies the same parametrized $SU(4)$ gates in a sequential descending pattern. (c)–(e) Task-specific readout protocols mapping quantum states to classical predictions. (c) Binary Classification Readout (used for Split CIFAR-100 and Quantum-Native tasks): The expectation value $\langle Z \rangle$ of the given qubit is measured and passed through a sigmoid activation to yield a binary class probability. (d) 10-class Classification Readout (used for Permuted MNIST): The probabilities of the first 10 computational basis states ($P(\dots 0000), \dots, P(\dots 1001)$) are extracted from the output distribution and normalized via a Softmax function to produce the categorical prediction. (e) Hybrid RL Readout (used for RL): The full probability distribution over all 2^N basis states serves as a dense feature vector. This vector is processed via a learnable linear projection (weighted sum) to generate the continuous policy actions and scalar value estimates, respectively.

ReLU activation, and a dense output layer with 10 units using a Softmax activation head. To test the impact of capacity, we vary the width of the hidden layer NW across $\{4, 6, 10, 16\}$ neurons.

- **Deep QNN:** We employ a VQC on $N = 10$ qubits (Fig. 5). The input is loaded via amplitude encoding. The variational *Ansatz* includes L layers (depth, DP) blocks. Each layer comprises generally two-qubit $SU(4)$ unitary gates applied to alternating pairs of neighboring qubits (even-odd pairs followed by odd-even pairs in a brickwall pattern). We vary the circuit depth $DP \in \{2, 4, 6, 9, 16\}$ to evaluate scalability.

- **QNN Readout:** A computational basis measurement is performed on all 10 qubits. The logarithmic probabilities of the first 10 bitstrings are interpreted as logits. These logits are processed by a softmax activation and fed into the cross-entropy loss.

To maximize local expressivity and avoid architectural bottlenecks, the fundamental building block of our variational *Ansatz* is the general parametrized two-qubit gate, $V \in SU(4)$. Unlike hardware-efficient *Ansätze* that rely on restricted gate sets (e.g., $R_y, R_z, CNOT$), a general two-qubit unitary possesses 15 degrees of freedom (ignoring global phase), allowing it to reach any state in the two-qubit Hilbert space.

We parametrize this gate using the canonical exponential map of the $\mathfrak{su}(4)$ Lie algebra. Let $\sigma_0 = I$ and $\{\sigma_1, \sigma_2, \sigma_3\} = \{X, Y, Z\}$ denote the standard Pauli basis. The gate $V(\theta)$ is defined by 15 trainable real parameters θ as [49]:

$$V(\theta) = \exp \left(-\frac{i}{2} \sum_{\substack{\alpha, \beta \in \{0,1,2,3\} \\ (\alpha, \beta) \neq (0,0)}} \theta_{\alpha, \beta} (\sigma_\alpha \otimes \sigma_\beta) \right). \quad (A1)$$

The summation runs over all 15 non-identity tensor products of Pauli matrices. By arranging these gates in a dense brickwall or ladder topology, the network creates a highly entangled state with a parameter manifold that is both compact and maximally expressive for its depth.

Training Protocol. All models are trained using the Adam optimizer. The batch size is 128 for all models.

- **Classical:** Learning rate 0.001, trained for 1 epoch per task.
- **Quantum:** Learning rate 0.03 trained for 5 epochs per task.

Plasticity Metrics. We also evaluate the average L_2 norm of the weight matrices (excluding biases for MLPs) and the average L_2 norm of the gradients computed on a fixed batch of data during training for each task.

2. Experiment 2: Supervised continual learning on split CIFAR-100

Task Formulation. We construct a sequence of $T = 3000$ binary classification tasks using the CIFAR-100 dataset. For each task, two distinct classes are sampled uniformly at random from the 100 available classes. The training and test datasets are filtered to include only samples from these two classes, and labels are remapped to $\{0, 1\}$.

Data Preprocessing. The $32 \times 32 \times 3$ RGB images are first converted to grayscale ($32 \times 32 \times 1$) and then flattened into vectors of dimension $d = 1024$ for both quantum and classical models. For both classical and quantum models, pixel values are L_2 -normalized to unit magnitude. This ensures that the classical model operates on the same geometric constraints as the quantum model for the input.

Model Architectures.

- **Classical MLP:** A deeper architecture is built for this visual task. It consists of an input layer ($d = 1024$), a first hidden layer with 256 ReLU neurons, a second hidden layer with varied ReLU neurons as $NW \in \{128, 256\}$, and a linear Dense output head for logits.
- **Deep QNN:** We use the same $N = 10$ qubit system as in experiment 1, but with significantly increased circuit depths to handle real-world image complexity. We test circuit depths $DP \in \{10, 20, 30\}$.
- **QNN Readout:** We measure the expectation value of the Pauli-Z operator on the last qubit as the logit, which is further used for classification loss.

Training Protocol. Both models are trained using the Adam optimizer with a fixed learning rate of $\eta = 0.001$ for 5 epochs per task. Batch size is 128.

Fisher Information Analysis. To quantify the effective learnability, we compute the trace of the FIM. To ensure temporal consistency, we define a fixed probe dataset D_{probe} comprising 64 randomly sampled images from the first task. The FIM trace is approximated as the expected squared L_2 norm of the score function (log-likelihood gradient) [53]:

$$\text{Tr}(F(\theta_t)) \approx \frac{1}{|D_{\text{probe}}|} \sum_{(\mathbf{x}, y) \in D_{\text{probe}}} \|\nabla_{\theta} \log p(y|\mathbf{x}; \theta_t)\|_2^2. \quad (\text{A2})$$

3. Experiment 3: Deep reinforcement learning

Environment. We employ the Ant-v4 environment from the MuJoCo physics engine via the Gymnasium library. The observation space is 27-dimensional, comprising joint angles, joint velocities, and center-of-mass coordinates (excluding external contact forces).

Reward Function. The agent is trained to maximize a composite reward signal R calculated at each time step:

$$R = \underbrace{v_x}_{\text{Forward}} + \underbrace{1.0}_{\text{Healthy}} - \underbrace{0.5\|\mathbf{a}\|^2}_{\text{Control Cost}} - \underbrace{5 \times 10^{-4}\|\mathbf{F}_{\text{clip}}\|^2}_{\text{Contact Cost}}, \quad (\text{A3})$$

where v_x is the forward velocity, 1.0 is a survival bonus if the ant does not fall, $\mathbf{a}$ is the action vector (control torque), and $\mathbf{F}_{\text{clip}}$ represents the external contact forces clipped to the range $[-1, 1]$. Note that while contact forces are excluded from the agent's input observation (to test plasticity on kinematic features), they are actively penalized in the reward to encourage efficient locomotion.

Data Preprocessing for Quantum Models. The 27-dimensional observation vector is zero-padded to 9-qubit dimension $d = 512$ (2^9) and L_2 -normalized to be compatible with quantum amplitude encoding.

Agent Architectures. We utilize the PPO algorithm with separate policy (π) and value (V) networks.

- **Classical Agent:** The policy network comprises two hidden dense layers with 1024 and 256 units (ReLU activation), followed by a linear action head of 8 outputs. The value network comprises two hidden layers with 256 units each.
- **Hybrid Quantum Agent:** The observation vector from the environment is first fed into a fixed-structure QNN acting on $N = 9$ qubits. The *Ansatz* consists of 12 layers of $SU(4)$ gates in ladder layout.
- **Quantum-Classical Interface:** To map the quantum state to classical policy features, we measure the probability distribution $p(\mathbf{z})$ of the basis states. These probabilities are processed by a scaled nonlinear activation function $h(\mathbf{z}) = \tanh(5.0 p(\mathbf{z}))$ to generate a dense feature vector. This vector is then projected via learnable linear heads (with no additional classical layers) to produce the stochastic action distribution vector and scalar value estimate.

Training Protocol. Agents are trained for a total of 30 million time steps using 16 parallel environments. We use a rollout buffer of 2048 steps, a batch size of 128, and 10 optimization epochs per update. The learning rate is set to 10^{-4} for the classical agent and 0.002 for the quantum agent.

4. Experiment 4: Continual learning on quantum data

Task Formulation. We generate a dataset of many-body quantum states derived from the one-dimensional Heisenberg XXZ Hamiltonian on a ring of $L = 10$ spins:

$$H = \sum_{i=1}^L (X_i X_{i+1} + Y_i Y_{i+1} + \Delta Z_i Z_{i+1}). \quad (\text{A4})$$

Periodic boundary conditions are assumed. We vary the anisotropy parameter Δ from -2.0 to 2 in steps of 0.02 . For each Δ , we compute the full eigenspectrum. We construct a sequence of 2000 tasks; each task involves binary classification between two distinct eigenstates $|\psi_i\rangle$ and $|\psi_j\rangle$ sampled uniformly from the spectrum. In other words, each task consists of 200 data points; each data point is a many-body wave function from either the i th (label 0) or j th (label 1) eigenstates of Hamiltonian with a certain Δ .

Input Representations.

- **Classical Input:** Since standard MLPs cannot process complex amplitudes directly, the state vector $|\psi\rangle \in \mathbb{C}^{1024}$ is decomposed. The real and imaginary parts are concatenated to form a real-valued input vector $\mathbf{x} \in \mathbb{R}^{2048}$.
- **Quantum Input:** The state $|\psi\rangle$ is loaded directly into the quantum register as the initial wavefunction, employing the native Hilbert space structure.

Model Architectures.

- **Classical MLP:** A feed-forward network with two hidden layers of NW neurons each (ReLU activation) and a single-unit output with sigmoid activation. We test $NW = 64, 256$.
- **Quantum Classifier:** A VQC with DP blocks of $SU(4)$ 2-qubit gates in a one-dimensional ladder pattern on 10 qubits. We test $DP = 6, 12$. The prediction logit is obtained by measuring the expectation value $\langle Z_0 \rangle$ of the first qubit, which further goes through the sigmoid activation in the binary cross-entropy loss.

Training Protocol. Both models are trained for five epochs per task with a batch size of 64 and a learning rate of 0.002 using the Adam optimizer.

APPENDIX B: ANALYTICAL DERIVATION OF FISHER INFORMATION SPECTRUM

In this section, we derive the asymptotic behavior of the FIM for classical and quantum neural networks trained with the binary cross-entropy (BCE) loss. We show that for classical networks (regardless of activation function), the unbounded growth of weight magnitudes leads to a collapse of the FIM trace, corresponding to a loss of plasticity. In contrast, we prove that the unitary nature of QNN guarantees a stable, non-vanishing FIM spectrum.

1. General form of FIM for diverse activation and loss functions

Consider a classification problem with input $\mathbf{x}$ and target y . Let the model produce a logit vector $\mathbf{f}(\mathbf{x}; \boldsymbol{\theta})$, which is transformed into a predictive distribution $p_j = \text{softmax}(\mathbf{f})_j$ (for multi-class) or $p = \sigma(f)$ (for binary).

By definition, the model Fisher Information Matrix is the expected outer product of the log-likelihood gradients: $F(\boldsymbol{\theta}) = \mathbb{E}_{\mathbf{x}, y \sim p_{\theta}} [\nabla_{\boldsymbol{\theta}} \mathcal{L} \nabla_{\boldsymbol{\theta}} \mathcal{L}^T]$. Applying the chain rule through the pre-activation logits $\mathbf{f}(\mathbf{x}; \boldsymbol{\theta})$ and exploiting the equivalence between the FIM and the generalized Gauss-Newton matrix, the expected second-order derivative of the model dynamics strictly vanishes under the model's own predictive distribution. This mathematically decouples the FIM into two distinct geometric components: the curvature of the output distribution and the Jacobian of the parameter manifold.

Consequently, the FIM reduces to the expectation over the data distribution:

$$F(\boldsymbol{\theta}) = \mathbb{E}_{\mathbf{x}} \left[\sum_{i,j} H_{ij} \nabla_{\boldsymbol{\theta}} f_i(\mathbf{x}) \nabla_{\boldsymbol{\theta}} f_j(\mathbf{x})^T \right], \quad (\text{B1})$$

where matrix $H_{ij} = \frac{\partial^2 \mathcal{L}}{\partial f_i \partial f_j}$. For the binary case (sigmoid), H reduces to the scalar $\xi = p(1-p)$. For the multi-class case (Softmax), the components are $H_{ij} = p_i \delta_{ij} - p_i p_j$,

where p_i is the probability of class i . In both cases, $\text{Tr}(F)$ represents the total sensitivity of the model's output distribution to parameter perturbations.

2. Plasticity collapse in classical networks

We analyze the behavior of Eq. (B1) under the condition of weight growth, a phenomenon ubiquitously observed in continual learning as shown in the main text.

Proposition 1. (Asymptotic Saturation). Consider a classical neural network $f(\mathbf{x}; \mathbf{w}) = \mathbf{w}_{\text{out}}^T \mathbf{h}(\mathbf{x}; \mathbf{w}_{\text{int}})$, where $\mathbf{w}_{\text{out}}$ are the readout weights and $\mathbf{h}$ is the latent space representation. If the magnitude of the weights grows unboundedly, i.e., $\|\mathbf{w}\| \rightarrow \infty$, then $\text{Tr}(F(\mathbf{w})) \rightarrow 0$.

Proof. In continual learning, to minimize loss on past data, the optimizer increases the magnitude of the weight vector. Let us parametrize this growth by a scaling factor λ as $\mathbf{w} \rightarrow \lambda \mathbf{w}_0$ when $\lambda \rightarrow \infty$.

Case 1: Sigmoid Binary Head. The logit output scales as $f_{\lambda}(\mathbf{x}) \approx \lambda f_0(\mathbf{x})$. The curvature term $\xi(\lambda) = \sigma(\lambda f_0)(1 - \sigma(\lambda f_0))$ decays exponentially: $\lim_{\lambda \rightarrow \infty} \xi(\lambda) \approx e^{-\lambda |f_0|}$. Since the gradients $\nabla_{\mathbf{w}} f$ scale at most polynomially with λ (e.g., linearly for ReLU or stay bounded for bounded activations), the product $\xi(\lambda) \|\nabla f\|^2$ vanishes as $\lambda \rightarrow \infty$.

Case 2: Softmax Multi-class Head. In classical classification architectures, regardless of the internal hidden activations (e.g., ReLU, Tanh, or periodic Sine), the final mapping to the unnormalized logits is a linear readout layer: $f_i(\mathbf{x}) = \mathbf{w}_{\text{out},i}^T \mathbf{h}(\mathbf{x})$, where $\mathbf{h}(\mathbf{x})$ is the representation from the penultimate layer. In continual learning, unbounded parameter growth implies $\|\mathbf{w}_{\text{out}}\| \rightarrow \infty$. Let us parametrize the divergent readout weights as $\mathbf{w}_{\text{out}}(\lambda) = \lambda \mathbf{w}_{\text{out},0}$ for $\lambda \rightarrow \infty$. Even if $\mathbf{h}(\mathbf{x})$ is strictly bounded by specific activation functions, the output logits scale linearly: $f_{\lambda,i}(\mathbf{x}) \approx \lambda f_{0,i}(\mathbf{x})$. For any given input $\mathbf{x}$, let $k = \text{argmax}_j f_{0,j}$ be the strictly dominant class, and let $\Delta_{\min} = \min_{j \neq k} (f_{0,k} - f_{0,j}) > 0$ be the minimum margin in the initial logit space. As $\lambda \rightarrow \infty$, the probability of the dominant class $p_k(\lambda)$ approaches 1, and any non-dominant class $j \neq k$ decays exponentially:

$$p_j(\lambda) = \frac{e^{\lambda f_{0,j}}}{\sum_i e^{\lambda f_{0,i}}} \leq e^{-\lambda(f_{0,k} - f_{0,j})} \leq e^{-\lambda \Delta_{\min}}. \quad (\text{B2})$$

Consequently, every element in the stiffness matrix $H_{ij}(\lambda) = p_i \delta_{ij} - p_i p_j$ becomes bounded by this exponential decay. Specifically, $H_{kk} = p_k(1 - p_k) \approx \mathcal{O}(e^{-\lambda \Delta_{\min}})$, and $H_{jj} \approx \mathcal{O}(e^{-\lambda \Delta_{\min}})$. Therefore, the matrix norm scales as:

$$\|H(\lambda)\| \sim \mathcal{O}(e^{-\lambda \Delta_{\min}}). \quad (\text{B3})$$

Crucially, the ultimate FIM trace is the product of this stiffness matrix H and the squared Jacobian $\|\nabla_{\mathbf{w}} \mathbf{f}\|^2$. By the chain rule, the gradient norm with respect to the network weights grows at most polynomially. For an

L -layer network, this growth is bounded by $\mathcal{O}(\lambda^{2L-2})$, regardless of whether the internal activations are unbounded (e.g., ReLU) or strictly bounded (e.g., Sine or Tanh), because the gradients are inevitably multiplied by the divergent subsequent weight matrices. Combining these two dynamics, the asymptotic limit of the FIM trace becomes a fundamental competition between polynomial gradient growth and exponential curvature collapse, leading to the vanishing FIM. This rigorously establishes that the geometric unboundedness of the Euclidean parameter space inevitably extinguishes the model's plasticity, a structural pathology that cannot be circumvented by simply tuning activation functions. ■

3. Preservation of plasticity in quantum neural networks

We now prove that QNNs are immune to this saturation mechanism due to the compactness of their parameter manifold.

Definition 1. (QNN Setup and Ansatz Structure). We define the QNN function $f(\mathbf{x}; \boldsymbol{\theta})$ as the expectation value of a measurement observable O (typically a Pauli string):

$$f(\mathbf{x}; \boldsymbol{\theta}) = s \langle 0 | U^\dagger(\mathbf{x}, \boldsymbol{\theta}) O U(\mathbf{x}, \boldsymbol{\theta}) | 0 \rangle. \quad (\text{B4})$$

Here, s is a fixed, non-trainable scalar. The unitary $U(\mathbf{x}, \boldsymbol{\theta})$ is composed of interleaved data encoding layers $S(\mathbf{x})$ and trainable layers $W(\boldsymbol{\theta})$. The trainable unitary is parameterized by rotation angles $\boldsymbol{\theta} = \{\theta_k\}_{k=1}^M$ generated by Pauli strings $H_k \in \{I, X, Y, Z\}^{\otimes N}$:

$$W(\boldsymbol{\theta}) = \prod_{k=1}^M e^{-i\theta_k H_k}. \quad (\text{B5})$$

Because $e^{-i\theta_k H_k}$ is periodic, the effective parameter space is the M -dimensional torus, $\mathbb{T}^M = [0, 2\pi)^M$.

Theorem 1. (Bounded FIM). For a QNN defined above, the trace of the FIM under BCE loss satisfies three conditions:

1. **Non-Vanishing:** The curvature term $\xi = p(1-p)$ is strictly lower-bounded.
2. **Non-Explosion:** The gradient norm is strictly upper-bounded, preventing chaotic landscapes.
3. **Recurrence:** The expected FIM trace over the parameter manifold is a non-zero constant.

Proof. Recall the FIM trace decomposition:

$$\text{Tr}(F) = \mathbb{E}_{\mathbf{x}} \left[\underbrace{p(1-p)}_{\xi} \times \underbrace{\sum_{k=1}^M (\partial_k f)^2}_{\|\nabla f\|^2} \right]$$

Part 1: Lower Bound on Curvature. In classical networks, $|f| \rightarrow \infty$ as weights grow. In QNNs, the output

is the expectation value of a Pauli observable O with spectral norm $\|O\|_\infty = 1$. Thus, the logit is strictly bounded for all $\boldsymbol{\theta}$:

$$|f(\mathbf{x}; \boldsymbol{\theta})| = |s \cdot \langle \psi | O | \psi \rangle| \leq |s| \cdot \|O\|_\infty = |s|. \quad (\text{B6})$$

Since $|f| \leq |s|$, the sigmoid derivative is strictly lower-bounded:

$$1/4 \geq \xi = \sigma(f)(1-\sigma(f)) \geq \sigma(|s|)(1-\sigma(|s|)) = C_{\min} > 0. \quad (\text{B7})$$

Part 2: Upper Bound on Gradients. In QNNs, we can bound the gradient using the parameter-shift rule. For Pauli generators with eigenvalues ± 1 , the gradient is exact:

$$\partial_k f = s \cdot \frac{1}{2} (\langle H_k \rangle_{\theta_k + \pi/2} - \langle H_k \rangle_{\theta_k - \pi/2}). \quad (\text{B8})$$

Since the expectation value $\langle H_k \rangle$ is bounded by $[-1, 1]$, the gradient is strictly bounded by the scaling factor:

$$|\partial_k f| \leq |s| \cdot \frac{1}{2} \cdot (1 + 1) = |s|. \quad (\text{B9})$$

Consequently, the total FIM trace has a strict global upper bound for any configuration of parameters:

$$\text{Tr}(F) \leq \underbrace{0.25}_{\max(\xi)} \sum_{k=1}^M (|s|)^2 = 0.25 \cdot M \cdot s^2 = C_{\max} < \infty. \quad (\text{B10})$$

This proves that the QNN landscape cannot become infinitely rugged or chaotic, i.e. extremely sensitive to training data.

Part 3: Stable Average Capacity. With the FIM strictly upper bounded, we finally show it does not statistically vanish with growing weights λ in CL. As parameters $\boldsymbol{\theta}$ traverse the compact torus $\mathbb{T}^M$, we examine the expected learning capacity. According to results from concentration of measure in quantum circuits, the expected squared gradient for a Pauli generator scales inversely with the effective dimension of the Hilbert space D_{eff} , where $D_{\text{eff}} \propto 2^N$ for sufficiently expressive *Ansätze* or global observables and $D_{\text{eff}} \propto \text{poly}(N)$ for barren plateau mitigated cases [58]:

$$\mathbb{E}_{\boldsymbol{\theta}}[(\partial_k f)^2] \propto \frac{1}{D_{\text{eff}}}. \quad (\text{B11})$$

Although this value depends on the system size or Hilbert space dimension, it is a static constant determined solely by the architecture and does not change over the training time. Unlike classical networks where the FIM decays exponentially with weight growth ($e^{-\lambda} \rightarrow 0$), the quantum gradient expectation remains invariant to the duration of training or λ . Combining parts 1, 2, and 3, we

establish that the QNN plasticity is permanently confined to a stable, active regime:

$$0 < \frac{C_{\min}}{D_{\text{eff}}} \leq \mathbb{E}[\text{Tr}(F)] \leq C_{\max} < \infty. \quad (\text{B12})$$

Note that the analytical derivation above specifically addresses binary classification tasks with the expectation value readout strategy, where the model output is defined as the expectation of a bounded observable. This formulation corresponds directly to the experimental setups employed in Split CIFAR-100 and quantum data classification.

4. Contrast with barren plateaus and dynamic plasticity

It is essential to distinguish the static challenge of barren plateaus (BP) from the dynamic loss of plasticity. BP is an initialization hurdle: as the number of qubits n scales, the variance of gradients [and thus $\text{Tr}(F)$] vanishes exponentially as $1/2^n$ for generic circuits and random initialization. However, QNNs utilizing amplitude encoding can represent high-dimensional features of length D using a logarithmic number of qubits $n = \log_2 D$. In our case, the 1024-dimensional visual features from CIFAR-100 require only 10 qubits ($2^{10} = 1024$). While the gradient variance in such circuits theoretically scales as $1/2^n = 1/D$, the dimensionality D in real-world data is not asymptotically exponentially large. This ensures that the system possesses a high information density that maintains expressive power while keeping the gradient signal at a manageable floor of $1/D$. Since D corresponds to fixed real-world input lengths, this gradient signal remains robustly within the trainable region despite the circuit's depth.

In contrast, loss of plasticity is a training-induced decay where the capacity collapses due to parameter drift into infinity. Our results show that for a QNN that has survived the BP regime (either through the use of high information density, local observables, specific circuit structures, or smart initialization), its learning capacity $1/D_{\text{eff}}$ remains a non-zero static constant. While the gradient signal may be small (scaling as $1/D$), it does not decay further as more tasks are learned.

To empirically verify this stability, we conducted a barren plateau analysis on the Split CIFAR-100 task setup. We varied the circuit depth from 1 to 60 and measured the squared norm of the gradient at initialization, averaged over 32 random seeds. As shown in Fig. 6, the gradient signal $\|\nabla_{\theta}\mathcal{L}\|^2$ remains robust and non-vanishing (averaging $\sim 5 \times 10^{-2}$) across all depths, including our standard study depth of 30. The model resides in a robustly trainable regime. This is fundamentally different from classical models, where the FIM trace decays exponentially toward zero over the task sequence, regardless of the initial capacity.

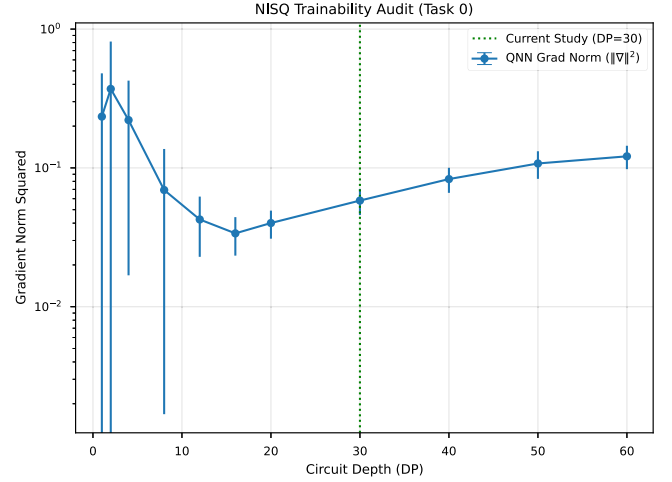


FIG. 6. Trainability of QNN for different depths. The squared norm of the gradient is plotted as a function of the variational circuit depth for the first task of the Split CIFAR-100 benchmark. The stable, non-vanishing gradient floor empirically confirms that the QNN operates in a trainable regime, ensuring that the tracked plasticity metrics are not artifacts of initial untrainability.

APPENDIX C: GENERALITY OF PLASTICITY PRESERVATION ACROSS QUANTUM CIRCUIT ANSTAZES

To assess the universality of our findings, we investigate whether the preservation of plasticity is specific to the general $SU(4)$ *Ansatz* employed in the main text or if it extends to standard, resource-constrained architectures common on near-term hardware.

We evaluate a hardware-efficient *Ansatz* (HEA) on the permuted MNIST benchmark. Unlike the main text model, which uses arbitrary two-qubit unitaries, this *Ansatz* enforces a structured topology with restricted gate sets, reducing the parameter density per layer.

The model operates on $N = 10$ qubits with a circuit depth of $D = 16$. Each layer consists of two sublayers:

- **Entanglement Layer:** Parametrized controlled-rotation- X gates ($CR_X(\theta)$) are applied in a ring connectivity. For qubits indexed $i \in \{0, \dots, N-1\}$, gates are applied between pairs $(i, (i+1) \pmod N)$.
- **Rotation Layer:** A sequence of local rotations is applied to every qubit individually: $R_y(\theta_1)R_z(\theta_2)R_y(\theta_3)$.

The similar structures are widely used in quantum machine learning due to their compatibility with near-term hardware qubit connectivity. The training protocol (learning rate, optimizer, task sequence) remains identical to the Permuted MNIST experiment described in the main text.

The results over a sequence of 500 tasks are visualized in Fig. 7. The HEA-based QNN maintains a stable test accuracy, with no observable downward trend. This result confirms that the loss of plasticity is not solved merely by



FIG. 7. Generality of plasticity preservation across quantum architectures. Test accuracy of a hardware-efficient *Ansatz* with 16 layers over a 500-task sequence of permuted MNIST. The architecture utilizes a ring topology of parametrized CR_x gates for entanglement, followed by a local $R_y - R_z - R_y$ sequence of rotations. Despite the restricted gate set compared to the $SU(4)$ model, the HEA maintains consistent learnability, confirming that the preservation of plasticity is a generic structural property of variational quantum circuits arising from their compact parameter space. The curve shows the moving average over 40-task windows, with shaded areas representing the standard deviation of the underlying data within each window.

the high expressivity of the $SU(4)$ gate but by the fundamental properties of the quantum optimization manifold. Even with a restricted gate set, the parameters remain confined to the compact torus, preventing the weight saturation observed in classical models. It also demonstrates that the plasticity advantage persists in architectures designed for physical hardware implementation, suggesting that current-generation quantum devices could potentially realize these benefits.

APPENDIX D: CONTROL EXPERIMENT WITH PERIODIC ACTIVATION FUNCTIONS

To rigorously test whether the preservation of plasticity in QNNs arises merely from the bounded nature of the quantum state output ($|\langle Z \rangle| \leq 1$) or fundamentally from the compactness of the parameter manifold, we perform a control experiment using a classical MLP with a periodic activation function.

We replace the standard ReLU activation with the sine function, $\phi(x) = \sin(x)$ [52]. Like the quantum expectation value, the output of a sine neuron is strictly bounded in $[-1, 1]$. If the boundedness of the representation is the sole factor for plasticity, this model should also exhibit plasticity.

The experiment is conducted on the Split CIFAR-100 benchmark (3000 tasks) using the same architecture as the standard MLP, but with the modifications that all hidden neurons use $\sin(x)$ activation. The epoch for training on each task is 20.

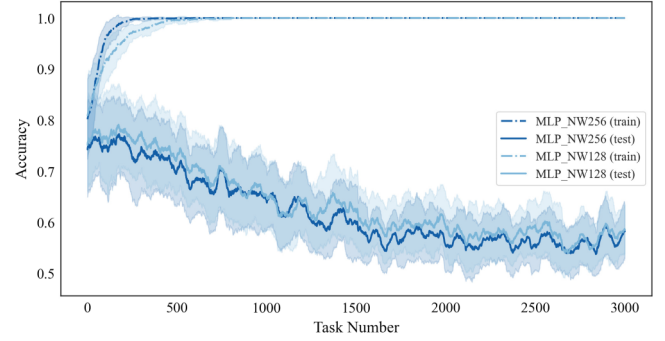


FIG. 8. Plasticity loss in classical MLPs with periodic activation on Split CIFAR-100. The plot shows the evolution of training accuracy (dashed lines) and test accuracy (solid lines) over 3000 binary classification tasks for network widths of 128 (light blue) and 256 (dark blue). Sin-MLPs maintain perfect training accuracy, while the test accuracy degrades continuously, exhibiting a wide generalization gap. This demonstrates that bounding the output representation in classical models is not sufficient to keep plasticity. All curves show the moving average over 50-task windows, with shaded areas representing the standard deviation of the underlying data within each window.

The results are presented in Fig. 8. We observe an evident decoupling between training and test performance and also a decay in test accuracy:

- **Training Accuracy (Dashed Lines):** The model maintains near-perfect training accuracy (100%) throughout the sequence. This indicates that the model retains high capacity.
- **Test Accuracy (Solid Lines):** Despite perfect training performance, the test accuracy degrades significantly over time, dropping from $\sim 75\%$ to $\sim 55\%$.

This behavior confirms the theoretical derivation before. Although the activation is bounded, the weights $\mathbf{w}$ reside in an unbounded Euclidean space. As the magnitude $\|\mathbf{w}\|$ increases to fit more tasks, the effective frequency of the sine neurons $[\sin(\mathbf{w}^T \mathbf{x})]$ increases. This transforms the optimization landscape into a high-frequency, extremely rugged terrain. While the optimizer can find a narrow minimum for the current training batch (hence 100% train accuracy), the resulting function generalizes poorly to the test set and becomes brittle to the distributional shifts inherent in continual learning.

Therefore, bounding the activation function is insufficient to preserve plasticity from the classical side. The indefinite plasticity observed in QNNs is therefore due to the compactness of the parameter manifold, which prevents the unbounded growth of complexity that drives the generalization collapse in periodic classical networks.

APPENDIX E: PERFORMANCE RESULTS WITH STANDARD DEVIATION ON SMOOTH WINDOW

In Figs. 2(b) and 4, we present the moving average of test accuracy to highlight the long-term plasticity trends.

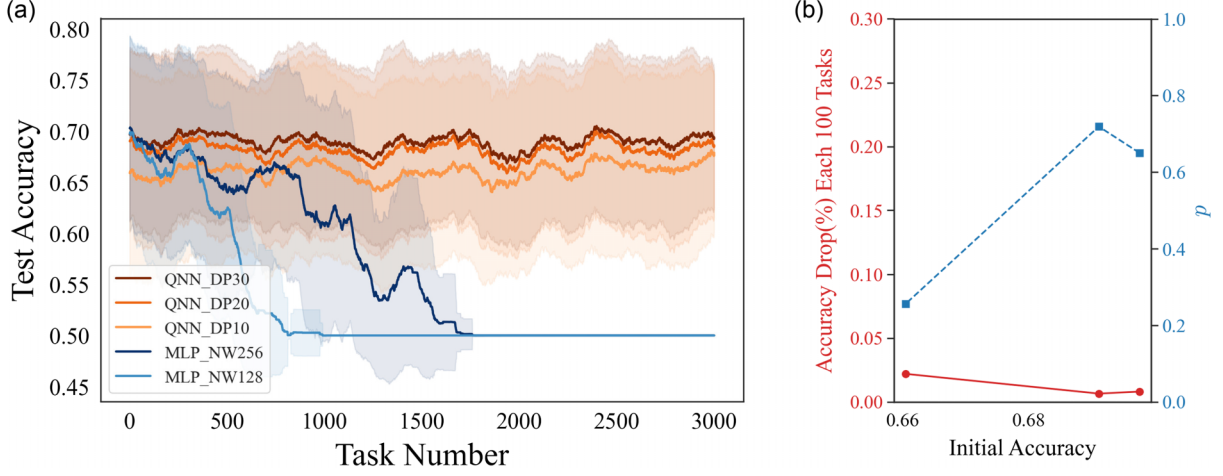


FIG. 9. Plasticity dynamics with performance variability on Split CIFAR-100. (a) Test accuracy as a function of task number for QNNs and MLPs, identical to Fig. 2(b) but visualizing the standard deviation (shaded regions) within each 150-task window. (b) The rate of plasticity loss (accuracy drop per 100 tasks) versus initial accuracy for QNN. Blue circles represent the corresponding p value for the linear estimation whose null hypothesis is zero slope.

To provide a complete picture of the learning stability and task-to-task variability, we present here the same results with full standard deviation bands included.

Figure 9 illustrates the results for the Split CIFAR-100 benchmark. The shaded regions represent the standard deviation of the test accuracy calculated within each 150-task moving window. The large band demonstrates the great variability of difficulty for image classification from different classes. Figure 9(b) directly shows the slope estimation of the QNN accuracy line. The estimated slope is negligible ($\approx 0.01\%$ drop per 100 tasks). Crucially, the p value associated with this slope estimation is $p \gg 0.05$. This statistical evidence confirms that the QNN performance is indistinguishable from a stationary process. The non-significant slope implies that the model is not slowly

dying, but rather fluctuating around a stable equilibrium of learnability, effectively preserving plasticity indefinitely.

Figure 10 presents the corresponding analysis for the quantum data continual learning task.

APPENDIX F: HYPERPARAMETERS IN ARCHITECTURES AND SETUPS

To ensure the reproducibility of our results and to demonstrate the universality of the plasticity preservation mechanism, we provide a detailed breakdown of the hyperparameters, network architectures, and initialization strategies employed across all four experimental domains: Permuted MNIST, Split CIFAR-100, Reinforcement Learning (Ant-v4), and Quantum Data Classification.

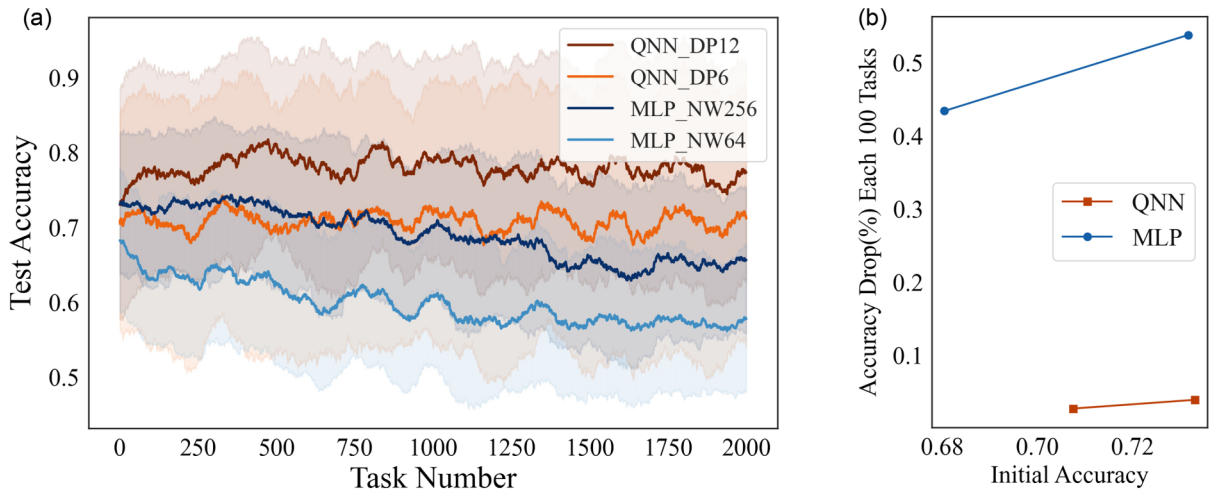


FIG. 10. Plasticity dynamics with performance variability on Quantum Data. (a) Test accuracy as a function of task number, identical to Fig. 4(b) but including the standard deviation (shaded regions) over 100-task moving windows. (b) Average accuracy drop per 100 tasks plotted against initial accuracy. The large value for MLPs confirms a consistent loss of plasticity, while QNNs maintain a near-zero drop rate, confirming the plasticity preservation.

TABLE I. Weight initialization strategies across experiments.

Experiment	Model	Weight initialization scheme
Permuted MNIST	Classical MLP	Glorot uniform
	Deep QNN	Uniform distribution $\mathcal{U}[0, 2\pi]$
Split CIFAR-100	Classical MLP	Glorot uniform
	Deep QNN	Uniform distribution $\mathcal{U}[0, 2\pi]$
RL Ant-v4	Classical agent	Orthogonal initialization (Gain = $\sqrt{2}$)
	Quantum agent	Standard normal distribution $\mathcal{N}(0, 1)$
Quantum data	Classical MLP	Glorot uniform
	Deep QNN	Uniform distribution $\mathcal{U}[-0.1, 0.1]$

TABLE II. Hyperparameter settings for the PPO agents in the Ant-v4 environment in RL task.

Hyperparameter	Classical and quantum agent
Horizon (steps per rollout)	2048
Mini-batch size	128
Number of Epochs (per update)	10
Discount factor	0.99
GAE parameter	0.95
Clip range (ϵ)	0.2
Value loss coefficient (c_1)	0.5
Entropy coefficient (c_2)	0.0
Max gradient Norm	0.5
Parallel environments	16
Total time steps	30,000,000

TABLE III. Comparison of model complexity in terms of total trainable parameters. Parameter counts for classical MLPs are derived from the specific architectures (hidden widths) described in the methods. For RL, the count includes both the policy and value networks. For QNNs, counts are based on the number of $SU(4)$ gates (15 parameters per gate) plus any classical trainable readout heads. QNNs achieve plasticity with orders of magnitude fewer parameters.

Experiment	Model	Configuration	Total parameters
Permuted MNIST	Classical MLP	Width 16 (1 hidden Layer)	$\approx 13,000$
	Deep QNN	Depth 16 (10 qubits)	$\approx \mathbf{2200}$
Split CIFAR-100	Classical MLP	Width 256 (2 hidden layers)	$\approx 330,000$
	Deep QNN	Depth 30 (10 qubits)	$\approx \mathbf{4100}$
RL Ant-v4	Classical agent	Policy + Value networks	$\approx 370,000$
	Hybrid quantum agent	QNN (12 Layers) + Linear heads	$\approx \mathbf{6200}$
Quantum data	Classical MLP	Width 256 (2 hidden layers)	$\approx 590,000$
	Deep QNN	Depth 12 (10 qubits)	$\approx \mathbf{1600}$

Crucially, these experiments span a wide range of configurations:

- **Model Complexity:** From compact QNNs (~ 200 parameters) to large classical MLPs ($> 500,000$ parameters).
- **Optimization Regimes:** From standard supervised learning with Adam to complex PPO updates.
- **Initialization Schemes:** From standard Glorot/Orthogonal initialization to simple uniform distributions.

Despite this heterogeneity in hyperparameters across tasks, the qualitative difference in learning dynamics, i.e.,

classical collapse versus quantum plasticity, remains consistent across all settings. This suggests that the observed advantage is not an artifact of specific tuning but arises from the fundamental structural differences.

The initial state of a network can significantly influence its training trajectory. As detailed in Table I, we employ industry-standard initialization schemes for the classical baselines to ensure they start from optimal conditions. Even with these stability-enhancing initializations, classical models fail to preserve plasticity over long task sequences. Conversely, QNNs utilize simple uniform or normal distributions for their rotation angles.

TABLE IV. Hyperparameter settings for the information plasticity analysis (model FIM analysis). These settings apply to the extended 3000-task investigation on Split CIFAR-100 benchmarking plasticity preservation mechanism.

Hyperparameter	Value
Analysis frequency	Every 5 tasks
Probe sample size (M)	64 (Task 0 training data)
Effective rank threshold	$\lambda_i > 10^{-4} \cdot \text{Tr}(F)$
Total CL tasks	3000
Optimizer	Adam (10^{-3})
Batch size	128

The specific hyperparameters for the RL agents are listed in Table II. To make a fair comparison, all the RL learning hyperparameters listed here are the same for classical and quantum agents.

Table III provides a comparative summary of the trainable parameter counts for all models. We note that QNNs maintain plasticity with orders of magnitude fewer parameters, highlighting the efficiency of QML. To track the underlying capacity dynamics, we perform a detailed model Fisher Information Matrix (FIM) analysis on the extended 3000-task Split CIFAR-100 sequence. The corresponding tracking hyperparameters, including analysis frequency and the effective rank threshold, are summarized in Table IV.

APPENDIX G: SUPPLEMENTARY BENCHMARKS AND ROBUSTNESS ANALYSIS

To quantify the robustness of the quantum plasticity advantage and isolate its structural origins, we conduct targeted comparative benchmarks.

L2 Regularization. We examine whether the loss of plasticity in classical networks is simply a consequence of unregularized weight growth. To test this, we evaluate a classical MLP with varying degrees of explicit weight decay on the Split CIFAR-100 task sequence. As shown in Fig. 11, L_2 regularization fails to preserve plasticity; the models still eventually collapse to the random guessing level (0.50). The effectiveness of this approach is highly sensitive to the choice of the weight decay hyperparameter. While small weight decays (10^{-4} and 10^{-3}) offer negligible help in maintaining accuracy over hundreds of tasks, a larger decay (10^{-2}) is overly restrictive, causing the network to fail to learn meaningful representations from the very beginning. The observations are consistent with Ref. [13]. This sensitivity suggests that L_2 regularization cannot fundamentally prevent the accumulation of dormant units or the decline of the representation's effective rank. These results highlight that extrinsic constraints in Euclidean space do not resolve the fundamental geometric saturation problem of the classical landscape.

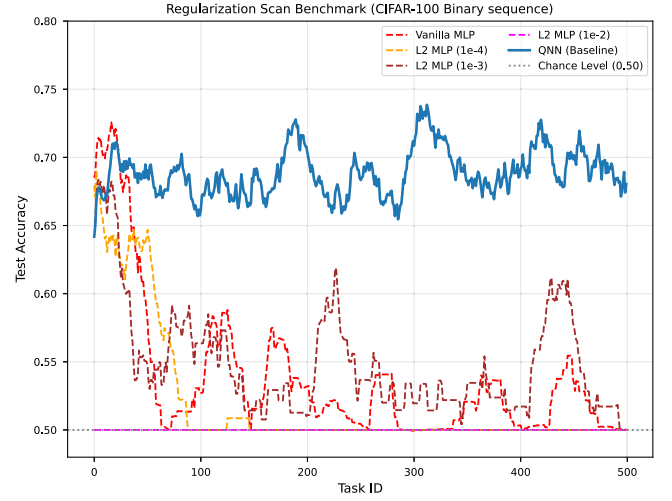


FIG. 11. Plasticity preservation under explicit regularization. Evolution of test accuracy on the CIFAR-100 Binary sequence benchmark across 500 tasks (Adam optimizer with learning rate 0.005, 5 epochs for one task). The performance of the baseline QNN is compared against a Vanilla MLP and classical MLPs with varying degrees of L_2 regularization. While the QNN sustains plasticity and maintains accuracy over time, explicit regularization fails to prevent catastrophic forgetting in the MLPs. Notably, small L_2 weight decays (10^{-4} and 10^{-3}) offer little help in preserving plasticity, with accuracy eventually degrading to the chance level (0.50). Conversely, a large L_2 weight decay (10^{-2}) is overly restrictive, causing the network to directly fail learning from the very beginning and remain at the chance level.

NISQ Noise Robustness. We verify the mechanism's robustness on hardware-realistic noisy intermediate-scale quantum devices by simulating a depolarizing noise model

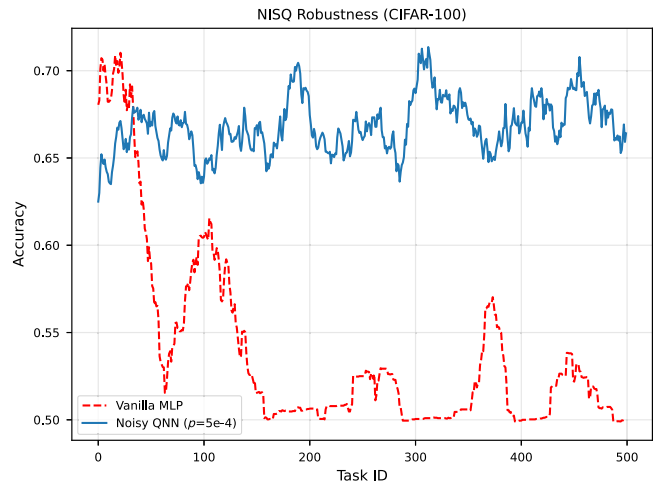


FIG. 12. Robustness to hardware noise in the NISQ era. Continual learning performance of the QNN under a depolarizing noise model (noise ratio $p_x = p_y = p_z = 5 \times 10^{-4}$, Adam optimizer with learning rate 0.005, 5 epochs for one task) on the Split CIFAR-100 benchmark. Despite the reduction in baseline performance, the QNN maintains its learnability over time, contrasting with the dynamic failure of classical architectures.

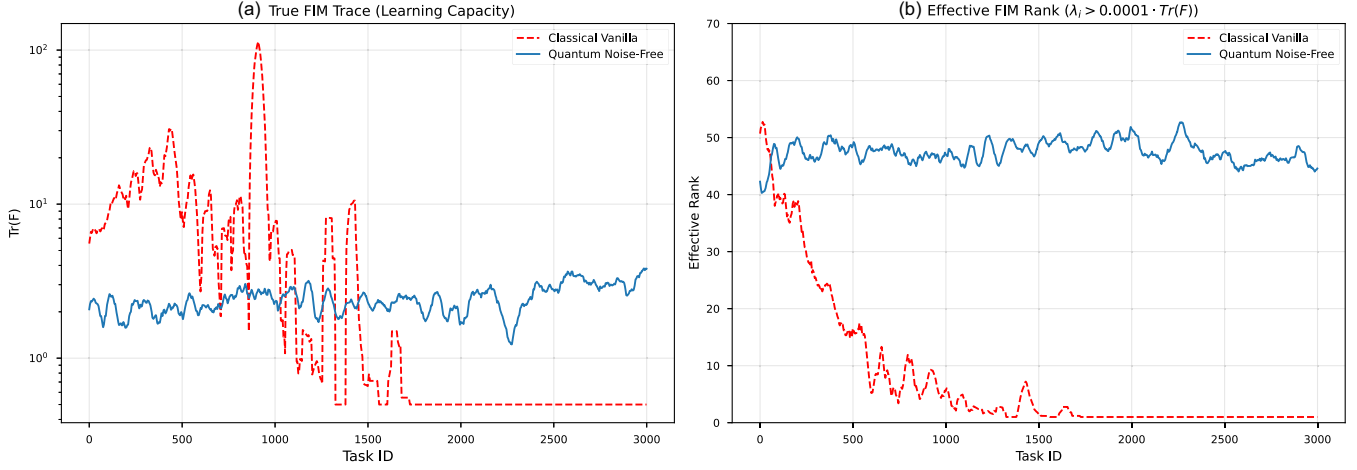


FIG. 13. Information-theoretic assessment of learning capacity. (a) Evolution of the trace of the model FIM over 3000 sequential tasks on the Split CIFAR-100 benchmark. (b) Monitoring of the effective FIM rank ($\lambda_i > 10^{-4} \text{Tr}(F)$). The functional dimensionality of the classical MLP (red dashed) collapses to unity, while the QNN (blue) preserves a high-rank landscape throughout the sequential learning process, providing a structural explanation for the persistent plasticity advantage.

($p_x = p_y = p_z = 5 \times 10^{-4}$). As visualized in Fig. 12, the QNN maintains a stable learning performance despite the systematic reduction in absolute baseline accuracy. This confirms that the geometric preservation of plasticity is robust to small quantum errors.

APPENDIX H: MODEL FISHER INFORMATION AND RANK ANALYSIS

To provide a more fundamental information-theoretic perspective on the loss of plasticity, we conduct an additional analysis using the *Model Fisher Information Matrix*, as opposed to the empirical FIM presented in the main text. Mathematically, while the empirical FIM is computed using the specific targets y_i present in the batch (representing the local curvature of the training loss surface), the model FIM is defined as the expectation over the model's own predictive distribution:

$$F = \mathbb{E}_{\mathbf{x} \sim p(\mathbf{x})} \mathbb{E}_{y \sim p(y|\mathbf{x}, \theta)} [\nabla_{\theta} \log p(y|\mathbf{x}, \theta) \nabla_{\theta} \log p(y|\mathbf{x}, \theta)^T]. \quad (\text{H1})$$

Physically, the model FIM represents the Riemannian metric of the model's manifold, measuring the sensitivity of the entire output distribution to parameter perturbations.

We observe that the model FIM and the empirical FIM exhibit qualitatively identical behavior in our lifelong learning experiments. This alignment stems from the fact that our models are consistently trained to a high-accuracy regime. When the model predictive distribution $p(y|\mathbf{x}, \theta)$ is well-optimized and highly confident (matching the true data distribution), the expectation over y in the model FIM calculation becomes dominated by the terms corresponding to the ground-truth labels. In this high-confidence regime, the information-theoretic curvature of the likelihood and the curvature of the actual loss surface

converge, rendering the empirical FIM a faithful and computationally efficient proxy for the fundamental learning capacity.

In our extended 3000-task investigation on Split CIFAR-100, we monitor both the trace of the model FIM and its *effective rank*, defined as the number of eigenvalues λ_i that satisfy $\lambda_i > 10^{-4} \cdot \text{Tr}(F)$. The analysis is conducted using a fixed probe set of $M = 64$ samples from the initial task, with the FIM being computed every 5 tasks. As shown in Fig. 13, the model FIM trace exhibits a qualitatively identical decay pattern to the empirical FIM in classical networks. More significantly, the effective rank of the classical MLP undergoes a catastrophic collapse, dropping from its initial value of 64 to exactly 1 by Task 1500. This indicates that the classical model's functional parameter space effectively implodes, leaving only a single effective direction for learning. In contrast, the QNN maintains a high and stable effective rank (averaging ~ 48) across the entire sequence, confirming that the unitary manifold preserves a high-dimensional landscape for continuous adaptation. These results support that the pattern of the trace and the effective rank of the FIM are qualitatively similar, and both are reliable proxies for the functional plasticity of the learner.

-
- [1] Y. LeCun, Y. Bengio, and G. Hinton, *Deep learning*, *Nature (London)* **521**, 436 (2015).
 - [2] M. I. Jordan and T. M. Mitchell, *Machine learning: Trends, perspectives, and prospects*, *Science* **349**, 255 (2015).
 - [3] R. Bommasani *et al.*, *On the opportunities and risks of foundation models*, [arXiv:2108.07258](https://arxiv.org/abs/2108.07258).
 - [4] S. Bilik *et al.*, *Towards phytoplankton parasite detection using autoencoders*, *Mach. Vis. Appl.* **34**, 101 (2023), <https://arxiv.org/abs/2303.08774>.

- [5] D. Guo *et al.*, *DeepSeek-R1 incentivizes reasoning in LLMs through reinforcement learning*, *Nature (London)* **645**, 633 (2025).
- [6] E. Strubell, A. Ganesh, and A. McCallum, *Energy and policy considerations for modern deep learning research*, in *Proceedings of the AAAI Conference on Artificial Intelligence* (2020), Vol. 34, p. 13693.
- [7] G. Ditzler, M. Roveri, C. Alippi, and R. Polikar, *Learning in nonstationary environments: A survey*, *IEEE Comput. Intell. Magazine* **10**, 12 (2015).
- [8] R. Hadsell, D. Rao, A. A. Rusu, and R. Pascanu, *Embracing change: Continual learning in deep neural networks*, *Trends Cognitive Sci.* **24**, 1028 (2020).
- [9] L. Wang, X. Zhang, H. Su, and J. Zhu, *A comprehensive survey of continual learning: Theory, method and application*, *IEEE Trans. Pattern Anal. Mach. Intell.* **46**, 5362 (2024).
- [10] I. J. Goodfellow, M. Mirza, D. Xiao, A. Courville, and Y. Bengio, *An empirical investigation of catastrophic forgetting in gradient-based neural networks*, [arXiv:1312.6211](https://arxiv.org/abs/1312.6211).
- [11] J. Kirkpatrick, R. Pascanu, N. Rabinowitz, J. Veness, G. Desjardins, A. A. Rusu, K. Milan, J. Quan, T. Ramalho, A. Grabska-Barwinska, D. Hassabis, C. Clopath, D. Kumaran, and R. Hadsell, *Overcoming catastrophic forgetting in neural networks*, *Proc. Natl. Acad. Sci. U.S.A.* **114**, 3521 (2017).
- [12] F. Zenke, B. Poole, and S. Ganguli, in *Proceedings of the 34th International Conference on Machine Learning, ICML'17* (JMLR.org, 2017), Vol. 70, pp. 3987–3995, <https://proceedings.mlr.press/v70/zenke17a.html>.
- [13] S. Dohare, J. F. Hernandez-Garcia, Q. Lan, P. Rahman, A. R. Mahmood, and R. S. Sutton, *Loss of plasticity in deep continual learning*, *Nature (London)* **632**, 768 (2024).
- [14] A. Chaudhry, P. K. Dokania, T. Ajanthan, and P. H. S. Torr, *Riemannian walk for incremental learning: Understanding forgetting and intransigence* (2018), pp. 556–572, [10.1007/978-3-030-01252-6_33](https://arxiv.org/abs/10.1007/978-3-030-01252-6_33).
- [15] J. T. Ash and R. P. Adams, in *Proceedings of the 34th International Conference on Neural Information Processing Systems, NeurIPS'20* (Curran Associates Inc., Red Hook, NY, 2020).
- [16] T. Berariu, W. Czarnecki, S. De, J. Bornschein, S. Smith, R. Pascanu, and C. Clopath, *A study on the plasticity of neural networks*, [arXiv:2106.00042](https://arxiv.org/abs/2106.00042).
- [17] E. Nikishin, M. Schwarzer, P. D'Oro, P.-L. Bacon, and A. Courville, in *Proceedings of the 39th International Conference on Machine Learning (PMLR, 2022)*, Vol. 162, pp. 16828–16847, <https://proceedings.mlr.press/v162/nikishin22a.html>.
- [18] Z. Abbas, R. Zhao, J. Modayil, A. White, and M. C. Machado, *Loss of plasticity in continual deep reinforcement learning*, [arXiv:2303.07507](https://arxiv.org/abs/2303.07507).
- [19] C. Lyle, Z. Zheng, E. Nikishin, B. A. Pires, R. Pascanu, and W. Dabney, in *Proceedings of the 40th International Conference on Machine Learning, ICML'23* (JMLR.org, 2023), <https://proceedings.mlr.press/v202/lyle23b.html>.
- [20] M. Elsayed and A. R. Mahmood, *Addressing loss of plasticity and catastrophic forgetting in continual learning*, [arXiv:2404.00781](https://arxiv.org/abs/2404.00781).
- [21] J. Biamonte, P. Wittek, N. Pancotti, P. Rebentrost, N. Wiebe, and S. Lloyd, *Quantum machine learning*, *Nature (London)* **549**, 195 (2017).
- [22] S. Lloyd, M. Mohseni, and P. Rebentrost, *Quantum algorithms for supervised and unsupervised machine learning*, [arXiv:1307.0411](https://arxiv.org/abs/1307.0411).
- [23] M. Schuld, I. Sinayskiy, and F. Petruccione, *The quest for a quantum neural network*, *Quantum Inf. Process.* **13**, 2567 (2014).
- [24] H.-Y. Huang, M. Broughton, M. Mohseni, R. Babbush, S. Boixo, H. Neven, and J. R. McClean, *Power of data in quantum machine learning*, *Nat. Commun.* **12**, 2631 (2021).
- [25] H.-Y. Huang, M. Broughton, J. Cotler, S. Chen, J. Li, M. Mohseni, H. Neven, R. Babbush, R. Kueng, J. Preskill, and J. R. McClean, *Quantum advantage in learning from experiments*, *Science* **376**, 1182 (2022).
- [26] M. Cerezo, G. Verdon, H.-Y. Huang, L. Cincio, and P. J. Coles, *Challenges and opportunities in quantum machine learning*, *Nat. Comput. Sci.* **2**, 567 (2022).
- [27] V. Havlíček, A. D. Córcoles, K. Temme, A. W. Harrow, A. Kandala, J. M. Chow, and J. M. Gambetta, *Supervised learning with quantum-enhanced feature spaces*, *Nature (London)* **567**, 209 (2019).
- [28] A. Pérez-Salinas, A. Cervera-Lierta, E. Gil-Fuster, and J. I. Latorre, *Data re-uploading for a universal quantum classifier*, *Quantum* **4**, 226 (2020).
- [29] J.-G. Liu and L. Wang, *Differentiable learning of quantum circuit born machines*, *Phys. Rev. A* **98**, 062324 (2018).
- [30] P.-L. Dallaire-Demers and N. Killoran, *Quantum generative adversarial networks*, *Phys. Rev. A* **98**, 012324 (2018).
- [31] B. Zhang, P. Xu, X. Chen, and Q. Zhuang, *Generative quantum machine learning via denoising diffusion probabilistic models*, *Phys. Rev. Lett.* **132**, 100602 (2024).
- [32] S.-X. Zhang, Z.-Q. Wan, C.-K. Lee, C.-Y. Hsieh, S. Zhang, and H. Yao, *Variational quantum-neural hybrid eigensolver*, *Phys. Rev. Lett.* **128**, 120502 (2022).
- [33] I. Cong, S. Choi, and M. D. Lukin, *Quantum convolutional neural networks*, *Nat. Phys.* **15**, 1273 (2019).
- [34] J. Li, X. Yang, X. Peng, and C.-P. Sun, *Hybrid quantum-classical approach to quantum optimal control*, *Phys. Rev. Lett.* **118**, 150503 (2017).
- [35] K. Beer, D. Bondarenko, T. Farrelly, T. J. Osborne, R. Salzmann, D. Scheiermann, and R. Wolf, *Training deep quantum neural networks*, *Nat. Commun.* **11**, 808 (2020).
- [36] H. Shen, P. Zhang, Y.-Z. You, and H. Zhai, *Information scrambling in quantum neural networks*, *Phys. Rev. Lett.* **124**, 200504 (2020).
- [37] L. Banchi, J. Pereira, and S. Pirandola, *Generalization in quantum machine learning: A quantum information standpoint*, *PRX Quantum* **2**, 040321 (2021).
- [38] W. Li and D.-L. Deng, *Recent advances for quantum classifiers*, *Sci. China Phys. Mech. Astron.* **65**, 220301 (2022).
- [39] P.-L. Zheng, J.-B. Wang, and Y. Zhang, *Efficient and quantum-adaptive machine learning with fermion neural networks*, *Phys. Rev. Appl.* **20**, 044002 (2023).
- [40] J. Jäger and R. V. Krems, *Universal expressiveness of variational quantum classifiers and quantum kernels for support vector machines*, *Nat. Commun.* **14**, 576 (2023).
- [41] J. Miao, C.-Y. Hsieh, and S.-X. Zhang, *Neural-network-encoded variational quantum algorithms*, *Phys. Rev. Appl.* **21**, 014053 (2024).

- [42] Y. Wu, J. Yao, P. Zhang, and X. Li, *Randomness-enhanced expressivity of quantum neural networks*, *Phys. Rev. Lett.* **132**, 010602 (2024).
- [43] W. Li, S.-X. Zhang, Z. Sheng, C. Gong, J. Chen, and Z. Shuai, *Quantum machine learning of molecular energies with hybrid quantum-neural wavefunction*, *Digit. Discov.* **4**, 2697 (2025).
- [44] Y.-Q. Chen and S.-X. Zhang, *Superior resilience to poisoning and amenability to unlearning in quantum machine learning*, [arXiv:2508.02422](https://arxiv.org/abs/2508.02422).
- [45] Z. Cui, P. Zhang, and Y. Tang, *Quantum flow matching*, [arXiv:2508.12413](https://arxiv.org/abs/2508.12413).
- [46] W. Jiang, Z. Lu, and D.-L. Deng, *Quantum continual learning overcoming catastrophic forgetting*, *Chin. Phys. Lett.* **39**, 050303 (2022).
- [47] C. Zhang *et al.*, *Quantum continual learning on a programmable superconducting processor*, *npj Quantum Inf.* **12**, 28 (2026).
- [48] K. He, X. Zhang, S. Ren, and J. Sun, in *Proceedings of the 2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)* (2016), pp. 770–778.
- [49] R. Dilip, Y.-J. Liu, A. Smith, and F. Pollmann, *Data compression for quantum machine learning*, *Phys. Rev. Res.* **4**, 043007 (2022).
- [50] K. Shen, B. Jobst, E. Shishenina, and F. Pollmann, *Classification of the fashion-MNIST dataset on a quantum computer*, [arXiv:2403.02405](https://arxiv.org/abs/2403.02405).
- [51] A. Kandala, A. Mezzacapo, K. Temme, M. Takita, M. Brink, J. M. Chow, and J. M. Gambetta, *Hardware-efficient variational quantum eigensolver for small molecules and quantum magnets*, *Nature (London)* **549**, 242 (2017).
- [52] V. Sitzmann, J. N. P. Martel, A. W. Bergman, D. B. Lindell, and G. Wetzstein, in *Proceedings of the 34th International Conference on Neural Information Processing Systems, NIPS '20* (Curran Associates Inc., Red Hook, NY, 2020).
- [53] A. Achille, M. Rovere, and S. Soatto, *Critical learning periods in deep neural networks*, [arXiv:1711.08856](https://arxiv.org/abs/1711.08856).
- [54] S. Kumar, H. Marklund, and B. V. Roy, *Maintaining plasticity in continual learning via regenerative regularization*, [arxiv:2308.11958](https://arxiv.org/abs/2308.11958).
- [55] E. Nikishin, J. Oh, G. Ostrovski, C. Lyle, R. Pascanu, W. Dabney, and A. Barreto, in *Proceedings of the 37th International Conference on Neural Information Processing Systems, NeurIPS '23* (Curran Associates Inc., Red Hook, NY, 2023).
- [56] S.-X. Zhang, J. Allcock, Z.-Q. Wan, S. Liu, J. Sun, H. Yu, X.-H. Yang, J. Qiu, Z. Ye, Y.-Q. Chen, C.-K. Lee, Y.-C. Zheng, S.-K. Jian, H. Yao, C.-Y. Hsieh, and S. Zhang, *Tensorcircuit: A quantum software framework for the NISQ era*, *Quantum* **7**, 912 (2023).
- [57] <https://github.com/sxzgroup/quantum-plasticity>.
- [58] J. R. McClean, S. Boixo, V. N. Smelyanskiy, R. Babbush, and H. Neven, *Barren plateaus in quantum neural network training landscapes*, *Nat. Commun.* **9**, 4812 (2018).