

Sharp-Wave Ripples Learning: A Bio-Inspired Incremental Learning Method

Yi Sun¹, Xiaochang Hu², Wenzhuo Zhang³, Zhiwei Wang³, Jiting Li^{1*}, Jian Li^{3*},
Xu Xin³ and Jia Jun¹

¹Academy of Military Sciences, China

²Beijing Institute of Basic Medical Sciences

³the College of Intelligence Science and Technology, National University of Defense Technology, China
lijiting_1993@126.com, lijian@nudt.edu.cn

Abstract

The primary challenge in incremental learning for deep neural networks (DNNs) lies in balancing memory stability with learning plasticity as they incrementally acquire new knowledge. Existing incremental learning methods typically require storing portions of old datasets or relying on specialized network architectures to achieve learning without forgetting. Recent neuroscience findings reveal that hippocampal sharp-wave ripples play critical roles in consolidating neocortex memory by internally triggering the reactivation of spontaneous neural circuits without external stimulus. Inspired by this biological mechanism, we propose an incremental learning method termed Sharp-Wave Ripples Learning (SWRL), a novel incremental learning framework composed of two functionally complementary modules: a neocortex-like network for long-term memory and a hippocampus-like model for rapid new-knowledge encoding. SWRL continuously integrates newly acquired knowledge from the hippocampus-like module into the neocortex-like network via a meta-weighted memory integration mechanism guided by internally generated SWR-like signals. Our framework is task-agnostic, architecture-agnostic, and exemplar-free. Experimental results on multiple public benchmarks demonstrate that SWRL achieves superior incremental learning performance under various scenarios across diverse DNN architectures, requiring no access to previous data or modifications to the base architecture. The code and appendix are both released at <https://github.com/SYVAE/SWRL>.

1 Introduction

Over the past decades, Deep Neural Networks (DNNs) have become fundamental technologies across various research domains such as autonomous driving, biomedical engineering [Notin *et al.*, 2024] and social simulation [Gao *et al.*, 2024a]. Despite existing remarkable achievements, DNNs currently exhibit limited life-long learning capabilities in incremental learning scenarios [Kudithipudi *et al.*, 2022].

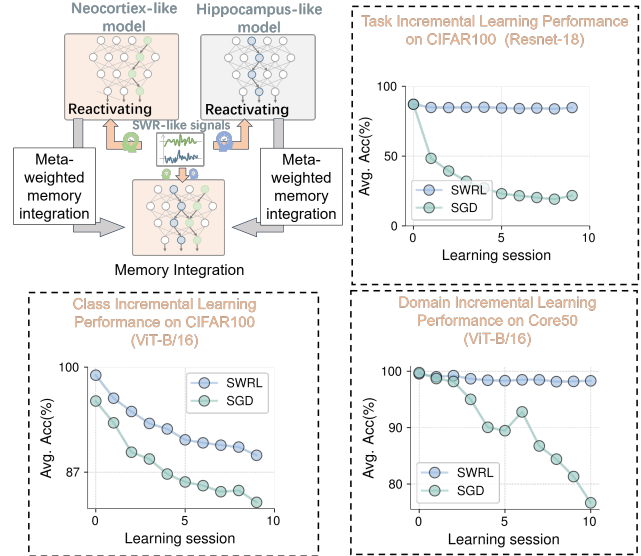


Figure 1: Superior performance and broad applicability of the proposed SWRL, which can be applied to diverse DNN architectures across three typical scenarios: task-incremental, class-incremental, and domain-incremental learning. The top-left section briefly illustrates the pipeline for meta-weighted memory integration and consolidation guided by SWR-like signals.

Specifically, training datasets are presented to DNNs sequentially in the incremental learning settings, demanding that DNNs dynamically adapt to non-stationary data streams while retaining previously acquired knowledge [Wang *et al.*, 2024]. However, existing commonly used optimization methods typically update synaptic weights greedily based on newly available data, resulting in severe performance degradation on prior tasks known as catastrophic forgetting (CF) [Wang *et al.*, 2024].

To tackle with the catastrophic forgetting problem, series of incremental learning methods have been proposed, including *replay-based* methods [Rebuffi *et al.*, 2017], which store and subsequently replay subsets of historical data during new learning sessions; *architecture-based* methods, which design specific architectures to minimize the interference between new and previously learned knowledge [Wang *et al.*, 2022b]; and *optimization-based* methods, which regulate the magni-

tude [Kao *et al.*, 2021] or direction [Zhang *et al.*, 2023] of weight updates. Nevertheless, when limiting the accessibility of old datasets on account of the learning efficiency, memory costs, privacy policies or other factors [Kudithipudi *et al.*, 2022], existing approaches struggle to balance the memory stability and learning plasticity of DNNs.

In contrast, the complementary learning system formed by the neocortex and hippocampus [Kudithipudi *et al.*, 2022] endows the brain with a superior capacity to learn continuously without forgetting. Recent advances in neuroscience reveal that memory integration and consolidation, within this system, are closely associated with hippocampal sharp-wave ripples [Norman *et al.*, 2019]. Specifically, after a learning episode concludes and the brain enters a wakeful rest state without external stimuli, bursts of hippocampal SWRs trigger spontaneous memory reactivations and synchronize the reinstatement of cortical neural circuits. During this process, neural circuits encoding long-term memories are consolidated, while new circuits are formed to integrate newly acquired knowledge from hippocampus.

Inspired by this biological mechanism, we propose an incremental learning method called Sharp-Wave Ripples Learning (SWRL) as illustrated in Figure 2. SWRL comprises two functionally distinct yet complementary modules: a neocortex-like model for storing long-term memory and a hippocampus-like model for rapidly acquiring new knowledge. The SWRL framework is designed to mimic the complementary interaction observed between the neocortex and hippocampus, involving two main stages: a fast learning stage and a memory integration stage. In the fast learning stage, the neocortex-like model is optimized to reduce the fusion gap between newly acquired and previously learned knowledge, serving as a "preview" to facilitate subsequent memory integration, while the hippocampus-like model is trained to fully learn new knowledge without concern of forgetting. Subsequently, during the memory integration stage, *SWRL generates SWR-like signals to reactivate memory-associated neural circuits, and updates the synaptic connections of the neocortex-like model to integrate newly learned knowledge in the hippocampus-like model, employing a novel meta-learning strategy.* In contrast to rehearsal exemplars of existing replay-based methods, *the SWR-like signals are generated internally by the neural network itself.* The contributions are summarized as follows:

- We propose a brain-inspired incremental learning method named SWRL, which simulates the collaborative mechanism between the hippocampus and neocortex. It achieves superior incremental learning performance without accessing prior datasets or modifying network architectures.
- We design a self-triggered meta-weighted memory integration algorithm within SWRL, that utilizes internally generated SWR-like signals to update neuron weights in a meta-weighted manner, effectively balancing memory stability and learning plasticity.
- SWRL is applicable to various incremental learning scenarios across diverse DNN architectures.

2 Related Work

2.1 Replay-based Incremental Learning

Replay-based methods alleviate CF by replaying datasets associated with previously learned knowledge during incremental learning [Rebuffi *et al.*, 2017; Chen *et al.*, 2024]. To enhance replay efficiency under constrained memory budgets, recent approaches leverage heuristic informations such as saliency information [Bellitto *et al.*, 2024] to select more representative rehearsal exemplars. Instead of storing raw training datasets, generative replay methods train generative models to synthesize pseudo-exemplars for rehearsal [Marañan *et al.*, 2021]. Advanced generative models, such as diffusion models [Kim *et al.*, 2024], have been employed to generate higher-fidelity rehearsal exemplars. Another complementary direction involves exploiting large-scale natural images to regularize representation drift via knowledge distillation [Feng *et al.*, 2023] or gradient de-conflicting strategies.

2.2 Optimization-based Incremental Learning

Optimization-based methods mainly include weight regularization, gradient projection, and sharpness-aware approaches. Weight regularization approaches stabilize old knowledge by explicitly restricting changes to weights that are important to previous tasks [Kirkpatrick *et al.*, 2017]. The importance of weights can be estimated using criteria such as gradient statistics or Fisher information [Kao *et al.*, 2021]. In contrast, gradient projection approaches stabilize memory by reprojecting gradients away from directions that harm previously learned knowledge [Lin *et al.*, 2022]. For instance, null-space gradient projection methods [Wang *et al.*, 2021] project gradients into the null space of previous feature space, effectively reducing hidden feature shifts. To generate better gradient subspaces that balance memory stability and learning plasticity of DNNs, recent methods have been developed to adopt adaptive orthogonal subspaces [Apolinario *et al.*, 2025] or synaptic meta-plasticity [Xie *et al.*, 2025]. Finally, sharpness-aware approaches aim to enhance the stability of DNNs by finding optimization solutions with flatter loss landscapes [Bian *et al.*, 2024], thereby explicitly reducing model sensitivity to weight perturbations.

2.3 Architecture-based Incremental Learning

Architecture-based methods require modifications to the network structure or depend on specific architectures. Early works, such as XdG [Masse *et al.*, 2018], introduced gating mechanisms to separate the networks into task-specific sub-networks using task embeddings, effectively alleviating inter-task interferences. Recently, with the rise of foundation models, MoE-style adapters have been widely employed to enhance the incremental learning performance of large language models or vision-language models [Yu *et al.*, 2024]. In contrast, model expansion approaches like NI-SSCL [Xie *et al.*, 2025], are proposed to learn both task-specific and task-shared modules through continuous genesis of model weights. Finally, prompt tuning approaches have become increasingly popular for transformer-based models. Instead of updating all backbone parameters, these approaches learn a small set of task-related prompt parameters to steer the frozen

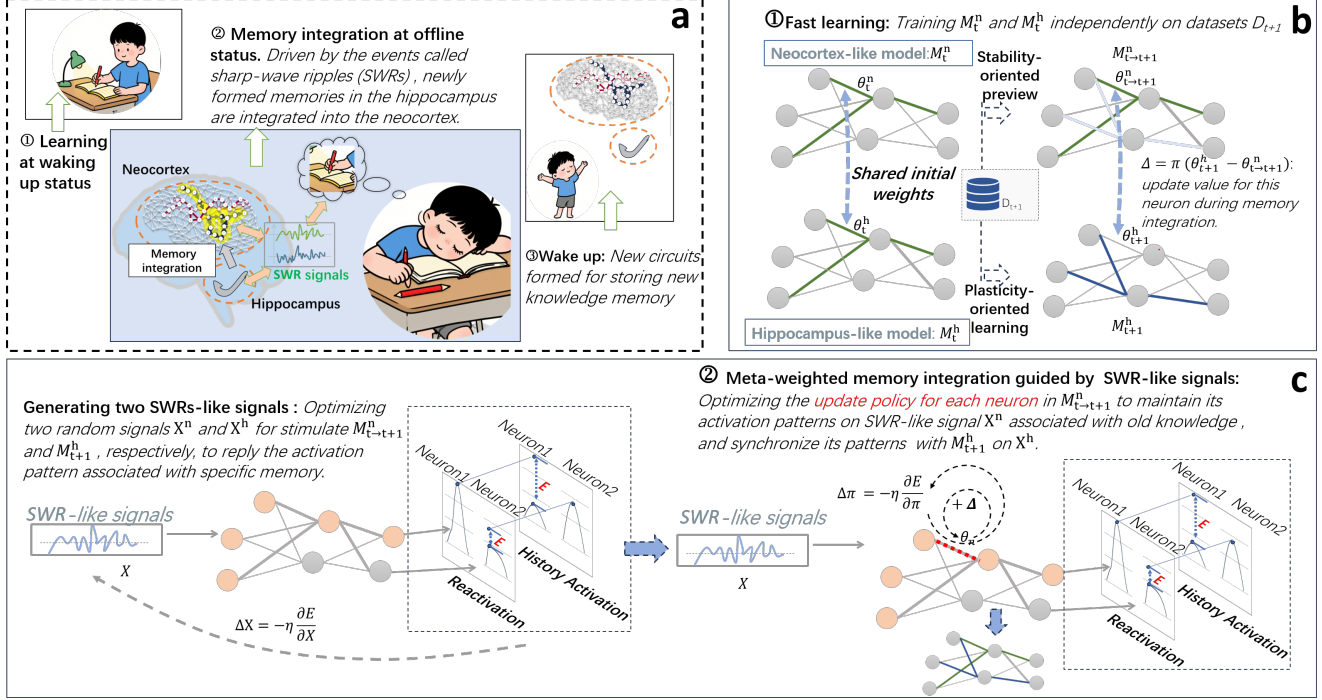


Figure 2: Overview of the proposed Sharp-Wave Ripples Learning (SWRL). (a) Biological mechanism. During offline brain states such as rest or sleep, newly formed memories in the hippocampus are gradually integrated into the neocortex. This process, known as systems memory consolidation, is facilitated by the synchronized reactivation of the hippocampus and neocortex, which is assumed to be triggered by events called sharp-wave ripples (SWRs). (b)-(c) Framework. The pipeline of the proposed SWRL method.

backbone for acquiring new knowledge [Wang *et al.*, 2022b]. Recent works have further improved prompt learning from different perspectives, such as tuning prompts along directions orthogonal to previously learned feature subspaces [Lu *et al.*, 2024], or combing prompt learning with replay-based methods [Ke *et al.*, 2025].

3 Method

In incremental learning, the training process of a DNN model is split into several sequential stages, where datasets in each stage contain specific learning objectives, such as adapting to new domains or acquiring task-specific skills. The proposed SWRL enables DNNs to learn without forgetting by mimicking the SWR-driven episodic replay phenomenon in the brain. As shown in Figure 2, SWRL mainly contains two stage, i.e., the fast learning stage and memory integration stage. Notably, it does not require access to prior datasets nor modifications to the model architecture.

3.1 Preliminary

Defining the sequential datasets as $\{\mathcal{D}_1, \dots, \mathcal{D}_{t+1}\}$ and the DNN model weights as Θ , the objective and optimization process at the $t + 1$ -th incremental learning stage are formulated

as:

$$\Theta_{t+1}^* = \arg \min_{\Theta_{t+1}} \underbrace{L_{t+1}(\Theta_{t+1} | \mathcal{D}_{t+1})}_{\text{Learning plasticity}} + \underbrace{\sum_{i=1}^t L_i(\Theta_{t+1} | \mathcal{D}_i)}_{\text{Memory stability}}, \quad (1)$$

where the model is supposed to achieve well performance both on the new dataset $\mathcal{D}_{t+1}$ and previous datasets after training, yet the optimization process can only access $\mathcal{D}_{t+1}$ as follows:

$$\Theta_{t+1} \leftarrow \Theta_{t+1} - \eta \nabla_{\Theta_{t+1}} L_{t+1}(\Theta_{t+1} | \mathcal{D}_{t+1}). \quad (2)$$

3.2 Fast Learning of Neocortex-like Model and Hippocampus-like Model

During the fast learning stage, a neocortex-like model $\mathcal{M}_t^n$ and a hippocampus-like model $\mathcal{M}_t^h$ are trained independently on the new dataset $\mathcal{D}_{t+1}$. $\mathcal{M}_t^n$ is trained to establish a more favorable neural circuitry foundation, aiming to reduce the integration gap between new and prior knowledge while prioritizing the maintenance of memory stability. In contrast, $\mathcal{M}_t^h$ is dedicated to acquiring new knowledge with a focus on learning plasticity.

Stability-oriented Preview Process of $\mathcal{M}_t^n$

Given a neocortex-like model $\mathcal{M}_t^n$ that have been trained on the previous t datasets, we employ a modified weighted null space projector [Xie *et al.*, 2025] to optimize $\mathcal{M}_t^n$. This

stability-oriented preview process aims to slightly enhance $\mathcal{M}_t^n$ on $\mathcal{D}_{t+1}$ while mitigating the forgetting of previously learned knowledge.

Consider a certain layer in $\mathcal{M}_t^n$ whose basic operator is a linear transformation, e.g., convolution layer and linear projection layer. We denote the weight matrix of this layer as $\theta_{t \rightarrow t+1}^n \in \mathbb{R}^{m \times n}$, initialized by θ_t^n . The update rule for it is formulated as:

$$\theta_{t \rightarrow t+1}^n \leftarrow \theta_{t \rightarrow t+1}^n - \eta \nabla_{\theta_{t \rightarrow t+1}^n} L_{t+1}(\theta_{t \rightarrow t+1}^n | \mathcal{D}_{t+1}) P_{null}, \quad (3)$$

where P_{null} is projection matrix designed to constrain the update, allowing learning on $\mathcal{D}_{t+1}$ while minimizing the change of prediction loss L_t on previous datasets $\mathcal{D}_t$. The calculation of P_{null} is detailed below.

Let the output of this layer as Z_t on $\mathcal{D}_t$, which can be calculated as:

$$Z_t = \theta_t^n X_t, \quad (4)$$

where $X_t \in \mathbb{R}^{n \times s}$ is the input feature for this layer corresponding to $\mathcal{D}_t$, and s denotes the number of spatial elements in the feature map. Then, the change in L_t , due to a weight update $\Delta\theta$, can be approximated via a first-order Taylor expansion:

$$\Delta L_t \approx \text{Tr}((\nabla_{Z_t} L_t)(\Delta Z_t)^T), \quad (5)$$

substituting $\Delta Z_t = \Delta\theta X_t$ yields:

$$\Delta L_t \approx \text{Tr}((\nabla_{Z_t} L_t)(\Delta\theta X_t)^T) = \text{Tr}(\Delta\theta X_t (\nabla_{Z_t} L_t)^T). \quad (6)$$

Let $M = X_t (\nabla_{Z_t} L_t)^T$, and performing singular value decomposition on M gives $M = U \Sigma V^T$. To ensure $\Delta L_t \approx 0$, $\Delta\theta$ should be constrained to the null space of M^T as expressed in Eq.(3). Then the corresponding projection matrix P_{null} is given as follows:

$$P_{null} = \hat{U} \hat{U}^T, \quad \hat{U} = [U_{n-k+1}, \dots, U_n], \quad (7)$$

where $\hat{U} \in \mathbb{R}^{n \times k}$ is formed by the last k left singular vectors $\{U_{n-k+1}, \dots, U_n\}$ of U , which correspond to the smallest k singular values in Σ . This yields a new neocortex-like model $\mathcal{M}_{t \rightarrow t+1}^n$ with weights $\theta_{t \rightarrow t+1}^n$, which establishes a more favorable neuron circuitry foundation for memory integration.

Plasticity-oriented Learning of $\mathcal{M}_t^h$

During this process, SWRL initializes the weights of the hippocampus-like model $\mathcal{M}_t^h$ with those of the neocortex-like model $\mathcal{M}_t^n$. It then employs common optimizers, such as SGD or Adam, to train $\mathcal{M}_t^h$ on the new datasets $\mathcal{D}_{t+1}$. The primary focus of SWRL in this plasticity-oriented learning stage is to update the network weights to minimize its loss on $\mathcal{D}_{t+1}$. Taking the weights θ^h of a certain layer as an example, this plasticity-oriented learning process is as follows:

$$\theta_{t+1}^h \leftarrow \theta_{t+1}^h - \eta \nabla_{\theta_{t+1}^h} L_{t+1}(\theta_{t+1}^h | \mathcal{D}_{t+1}), \quad (8)$$

yielding a new hippocampus-like model $\mathcal{M}_{t+1}^h$ with weights θ_{t+1}^h which encodes more comprehensive new knowledge.

3.3 Meta-weighted Memory Integration Guided by SWR-like Signals

The fast learning stage yields a new neocortex-like model $\mathcal{M}_{t \rightarrow t+1}^n$ and a hippocampus-like model $\mathcal{M}_{t+1}^h$. To integrate the knowledge of $\mathcal{M}_{t+1}^h$ into $\mathcal{M}_{t \rightarrow t+1}^n$, SWRL employs a novel meta-weighted memory integration algorithm, guided by SWR-like signals, to update $\mathcal{M}_{t \rightarrow t+1}^n$.

Mathematical Formulation of Memory Integration

For brief illustration, we take a certain layer as an example to introduce the memory integration. The integration process for the weight matrix θ can be initially formulated as:

$$\theta_{fusion} = \theta_{t \rightarrow t+1}^n + \underbrace{\pi \mathbf{1}_n^T \odot (\theta_{t+1}^h - \theta_{t \rightarrow t+1}^n)}_{\Delta\theta}, \quad \pi \in \mathbb{R}^{m \times 1}, \quad (9)$$

where $\odot$ denotes the Hadamard product, and $\mathbf{1}_n$ represents an n -dimensional all-ones column vector. The proposed SWRL will adaptively access the importance (π) of neurons to different memory, and accordingly adjust the update policy for each neuron by optimizing π for balancing learning plasticity and memory stability.

Meta-weighted Memory Integration by SWR-like Signals

Neuron circuit reactivations in the neocortex are closely coordinated with hippocampal sharp-wave ripples (SWRs) [Norman *et al.*, 2019]. This cortico-hippocampal dialogue, is believed to be crucial for integrating newly acquired information from the hippocampus into long-term storage within the neocortex [Kudithipudi *et al.*, 2022]. Inspired by this mechanism, we design a meta-weighted integration policy guided by SWR-like signals.

Specifically, SWR-like signals are tensors generated by the DNN models themselves, which can reactivate the models to reproduce the neuron activation patterns associated with specific memory of knowledge without accessing the original training datasets. We denote the signals that reactivate neuron circuits encoding previous knowledge as X^n , and those for newly learned knowledge as X^h . Details of generating SWR-like signals are provided in the next subsection.

As shown in Eq.(10), SWRL optimizes π to preserve the activation patterns of $\mathcal{M}_{t \rightarrow t+1}^n$ on X^n , while synchronizing its activation patterns with $\mathcal{M}_{t+1}^h$ on X^h . The optimization of π is a bi-level meta learning process, consisting of an inner optimization (InnerOpt) and an outer optimization (OuterOpt), described as follows:

$$\begin{aligned} \pi^* = & \underbrace{\arg \max_{\pi} \sum_{i \in \{h, n\}} \mathcal{C}_{syn}(\mathcal{M}_{t \rightarrow t+1}^n(X^i; \underbrace{\theta_{t \rightarrow t+1}^n + \Delta\theta}_{\text{InnerOpt}}, a^i)}_{\text{OuterOpt}}, \\ a^n = & \mathcal{M}_{t \rightarrow t+1}^n(X^n; \theta_{t \rightarrow t+1}^n), \quad a^h = \mathcal{M}_{t+1}^h(X^h; \theta_{t+1}^h), \end{aligned} \quad (10)$$

where $\mathcal{C}_{syn}(a, b)$ measures the alignment of activation patterns using cosine similarity. SWRL adaptively evaluates the sensitivity of different knowledge to changes in neuron

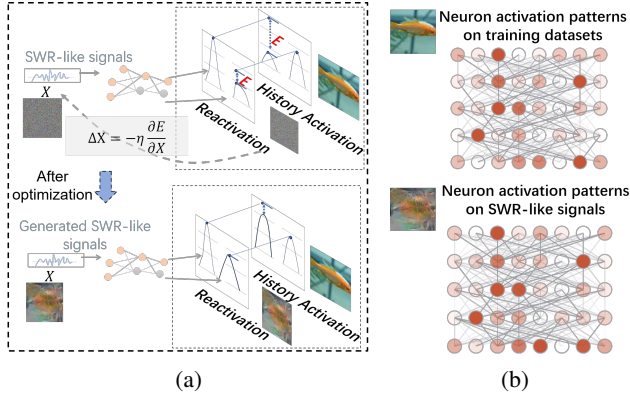


Figure 3: Generating SWR-like signals using neuron activation-consistency prior. (a) Signal generation. (b) Model reactivation. Each circle represents a real neuron in Resnet-18 trained on CIFAR100, and the darker colors indicate higher levels of activation.

weights, and accordingly adjust the fusion weights for each neuron. The fusion results are given by:

$$\theta_{fusion}^* = \theta_{t \rightarrow t+1}^n + \pi^* \mathbf{1}_n^T \odot (\theta_{t+1}^h - \theta_{t \rightarrow t+1}^n). \quad (11)$$

Generating SWR-like Signals Using Activation Consistency

The SWR-like signals $\{X^n, X^h\}$ are applied to the evaluation of forgetting and synchronization level during memory integration. The proposed SWRL optimizes randomly initialized tensors X to obtain such signals, each associated with specific knowledge. Taking the generation of X^h for $\mathcal{M}_{t+1}^h$ as an example, and assuming that produce SWR-like signals that reactivate the model for recalling newly learned knowledge about "what is y_{t+1} ?", then this is formulated as:

$$X^h = \arg \min_X \sum_{i=1}^N \left\{ (\mathcal{M}_{t+1}^{h,i}(X) - \mu_i)^2 + (1 - C_{syn}(\mathcal{M}_{t+1}^{h,i}(X), \mu_i)) \right\} + L_{t+1}(\mathcal{M}_{t+1}^h(X), y_{t+1}), \quad (12)$$

where $\mathcal{M}_{t+1}^{h,i}(X)$ denotes the activation of the i -th hidden (convolutional or linear projection) layer given input X , and $\{\mu_i\}$ represents the recorded historical activation levels of this layer. As shown in Figure 3, to generate signals that reactivate the circuits encoding knowledge of a "Gold fish" in ResNet-18, μ_i is the running mean of the corresponding layer, which can be obtained from its cascaded batch normalization layer. Here, y_{t+1} corresponds to the prediction logit vector in which the entry associated with "Gold fish" has the highest value. As shown in Figure 3(b), the activation patterns elicited by the generated signals closely resemble those produced by the original relevant inputs, demonstrating successful knowledge recall.

4 Experiments

We compare the proposed SWRL with state-of-the-art methods under three typical incremental learning settings, includ-

Algorithm 1 SWRL

Require: Initial neocortex-like model: $\mathcal{M}_1^n$, training dataset: $\{\mathcal{D}_t\}_{t \leq T}$. The number of incremental learning stage is T .

Ensure: neocortex-like model: $\mathcal{M}_T^n$

- 1: Generating SWR-like signals of learned knowledge according to Eq. (12): (X_1^n)
- 2: **for** $t \in [2, T]$ **do**
- 3: ▼ *Fast learning stage*
- 4: Stability-oriented preview process of neocortex-like model $\mathcal{M}_{t-1}^n: \mathcal{M}_{t-1}^n \xrightarrow{\text{Eq. (3)}} \mathcal{M}_{t-1 \rightarrow t}^n$
- 5: Initializing hippocampus-like model $\mathcal{M}_t^h$ by $\mathcal{M}_{t-1 \rightarrow t}^n$
- 6: Plasticity-oriented learning of hippocampus-like model $\mathcal{M}_t^h: \mathcal{M}_t^h \xrightarrow{\text{Eq. (8)}} \mathcal{M}_t^h$
- 7: ▼ *Memory integration stage*
- 8: Generating SWR-like signals on datasets $\mathcal{D}_t$ by Eq.(12)
- 9: Meta-weighted memory integration guided by SWR-like signals by Eq.(10): π^*
- 10: Calculating new neocortex-like model $\mathcal{M}_t^n$ by Eq.(11)
- 11: **end for**

ing task incremental learning, class incremental learning, and domain incremental learning [van de Ven *et al.*, 2022].

Datasets

This evaluation utilizes datasets from vision tasks, and employs DNN architectures including various convolutional neural networks and transformers. The employed datasets are described as follows:

- In task-incremental scenarios, we incrementally train Resnet-18 [He *et al.*, 2016] on CIFAR100 [Krizhevsky *et al.*, 2009] and TinyImageNet [Wu *et al.*, 2017]. Following the experimental setup in [Lin *et al.*, 2022], CIFAR100 dataset is split into 10 non-overlapping subsets, resulting in the task incremental learning dataset CIFAR100-TIL-s10. TinyImageNet is similarly divided into 25 non-overlapping subsets, forming TinyImageNet-TIL-s25.
- In class incremental learning, we incrementally train ViT-B/16 [He *et al.*, 2016] on CIFAR100 [Krizhevsky *et al.*, 2009]. Following the settings in [Gao *et al.*, 2024b], CIFAR100 dataset is split into 10 non-overlapping subsets, resulting in the class incremental learning dataset CIFAR100-CIL-s10. Different from the setup in task-incremental scenarios, there is no task identity provided in this setting, which makes the class incremental learning is more challenging than task incremental learning.
- In domain-incremental scenarios, we follow the experimental settings of [Shi and Wang, 2024], where Sequential CORE50 [Lomonaco and Maltoni, 2017] is used to construct domain incremental learning dataset. CORE50 is a real-world incremental dataset that includes 50 do-

Methods	CIFAR100-TIL-s10		TinyImageNet-TIL-s25	
	Avg. Acc(%) $\uparrow$	BWT(%) $\uparrow$	Avg. Acc(%) $\uparrow$	BWT(%) $\uparrow$
Adam-NSCL	73.77	-1.6	60.31	-6.05
Connector	79.79	-0.92	64.61	-6
GPCNS	74.40	-2.16	-	-
NI-SSCL	75.51	-0.76	-	-
HINT	77.46	-	66.10	-
BECAME	81.66	-0.80	66.49	-4.97
SWRL(ours)	83.35	-0.76	68.52	-7.48

Table 1: Comparison results on task incremental datasets CIFAR100-TIL-s10 and TinyImageNet-TIL-s25 (Resnet-18).

mestic objects collected from 11 domains (image size: 128×128 ; total of 120,000 images).

Evaluation Metrics

We use the average classification accuracy (Avg. Acc) to evaluate the overall incremental learning performance. Let $A_{m,t}$ represents the classification accuracy of the neural network on the t -th datasets after incrementally learning on m datasets, and let K be the total number of learning sessions. The Avg. Acc can be calculated as follows:

$$\text{Avg. Acc} = \frac{1}{K} \sum_{t=1}^K A_{K,t}. \quad (13)$$

In addition, we employ the backward transfer metric (BWT) and forgetting measure (FM) [Wang *et al.*, 2024] to evaluate the degree of forgetting during the incremental learning process. BWT and FM can be calculated as:

$$\text{BWT} = \frac{1}{K-1} \sum_{t=1}^{K-1} (A_{K,t} - A_{t,t}), \quad (14)$$

$$\text{FM} = \frac{1}{K-1} \sum_{t=1}^{K-1} f_{K,t}, \quad f_{K,t} = \max_{i \in [1, K-1]} (A_{i,t} - A_{K,t}). \quad (15)$$

4.1 Comparison of SWRL with Existing Works

The proposed SWRL achieves state-of-the-art performance across various incremental learning scenarios. Specifically, in the task-incremental learning setting, we compare SWRL with Adam-NSCL [Wang *et al.*, 2021], Connector [Lin *et al.*, 2022], GPCNS [Yang *et al.*, 2024], NI-SSCL [Xie *et al.*, 2025], HINT [Krukowski *et al.*, 2025], and BECAME [Li *et al.*, 2025]. SWRL attains a remarkable average accuracy of 83.35% on CIFAR100-TIL-s10, while exhibiting minimal forgetting (as reflected in its superior BWT score -0.76%), thereby outperforming existing methods. Moreover, on the more challenging TinyImageNet-TIL-s25 dataset, although its BWT is not the highest, it achieves the best overall performance with average accuracy of approximately 68.52%, substantially outperforming other approaches. Additional results are reported in Appendix, where SWRL achieves 94.12% average accuracy using ViT-B/16 on TinyImageNet-TIL-s25.

Class incremental learning is more challenging than task incremental learning, as task identities are not provided. In our CIFAR100 class-incremental learning experiment, we employed a commonly used Vision Transformer, ViT-B/16

Methods	Avg. Acc(%) $\uparrow$	FM(%) $\downarrow$
L2P [Wang <i>et al.</i> , 2022c]	86.38	5.88
DualPrompt [Wang <i>et al.</i> , 2022b]	86.61	5.86
C-prompt [Gao <i>et al.</i> , 2024b]	87.82	5.06
InfLoRA [Liang and Li, 2024]	88.31	-
PEGP [Qiao <i>et al.</i> , 2025] $\dagger$	85.43	6.38
DMNSP [Kang <i>et al.</i> , 2025] $\dagger$	87.59	-
SWRL(ours)	89.15	4.47

Table 2: Comparison results on class incremental datasets CIFAR100-CIL-s10 (ViT-B/16). $\dagger$: best results reported in the reference paper.

Methods	Buffer size	Avg. Acc(%) $\uparrow$
L2P [Wang <i>et al.</i> , 2022c]	50/class	81.07
C-Prompt [Wang <i>et al.</i> , 2022a]		82.86
L2P [Wang <i>et al.</i> , 2022c]		78.33
S-Prompts [Wang <i>et al.</i> , 2022a]	0/class	83.13
C-Prompt [Gao <i>et al.</i> , 2024b]		85.31
PINA-D+SOYO [Wang <i>et al.</i> , 2025]		90.77
DUCT [Zhou <i>et al.</i> , 2025]		94.47
SWRL(ours)		98.29

Table 3: Comparison results on domain incremental datasets Core50 (ViT-B/16). "Buffer size" indicates the number of samples retained from previous domains.

[Dosovitskiy *et al.*, 2021], as the base model and applied SWRL for training. The results in Table 2 show that SWRL not only enhances the performance of convolutional networks in task incremental learning but is also effective for transformer-based models, even in the more challenging class incremental setting. Compared with state-of-the-art methods such as C-prompt [Gao *et al.*, 2024b] and DMNSP [Kang *et al.*, 2025], SWRL achieves a better balance between stability and plasticity, attaining an overall accuracy of 89.15%. Crucially, owing to its architecture-agnostic and task-agnostic design, SWRL is more generalizable than existing leading prompt-tuning methods, which are often effective only for transformer-based models.

In domain-incremental learning settings, models are required to continuously adapt to new data distributions. A Vision Transformer ViT-B/16 pre-trained on ImageNet-21K is used as the base model. The results, presented in Table 3, show that SWRL achieves an overall accuracy of 98.29%, outperforming both replay-based methods using additional data buffers and exemplar-free prompt-based methods such as DUCT [Zhou *et al.*, 2025]. In summary, the results in Tables 1 to 3 demonstrate that SWRL is a general incremental learning method applicable to various incremental learning scenarios across diverse deep neural network architectures.

4.2 Analysis of SWRL

Minimization of the Integration Gap

The designed fast learning process, consisting of stability-oriented preview process of $\mathcal{M}_t^n$ and plasticity-oriented learning of $\mathcal{M}_t^h$, effectively closes the integration gap between the new-learned memory and previous long-term memory in $\mathcal{M}_t^n$. We visualize the loss landscape of models on the 1st and 2nd task in Figure 4 by training Resnet-18 on CIFAR100-TIL-s10. It can be seen that the preview process

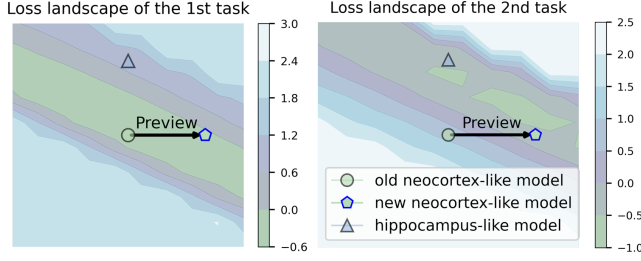


Figure 4: Landscape of models on different tasks.

optimize $\mathcal{M}_t^n$ to move to the low loss region of the 2nd task, while keeping its location at the low loss region of the 1st task. Results reported in Table 4 verify the effectiveness of the designed fast learning process.

Balancing Memory Stability and Learning Plasticity

The neuron circuits of the neocortex-like model will be changed when integrating new learned knowledge from the hippocampus-like model, inevitably disturbing the memory stability. To balance the learning plasticity and memory stability, the proposed SWRL adaptively optimizes the integration policy, as described in Eq.(10).

Methods	preview	meta	Avg. Acc(%) $\uparrow$	BWT(%) $\uparrow$
SWRL(ours)	✗	✓	64.75	-3.7
	✓	✗	80.65	-2.3
	✓	✓	83.35	-0.76

Table 4: Comparison results on CIFAR100-TIL-s10.

We take memory integration comparison experiments on Resnet-18 using CIFAR100-TIL-s10. Specifically, we compare the performance when using meta-weighted memory integration guided by SWR-like signals, and adopting linear fusion with constant fusion coefficient matrices, respectively. As illustrated in Table 4, the designed memory integration mechanism effectively enhances the overall accuracy while alleviating the forgetting. Figure 5 further presents the comparison results on ResNet-18 using TinyImageNet-TIL-s25, showing that the meta-weighted memory integration achieves better incremental learning performance.

For more intuitively demonstrating the changes of neocortex-like model after memory integration, we further visualize the neuron-circuits activation pattern in Figure 6. For clarity, we highlight partial neurons with higher activation levels. Figure 6 demonstrates that, after memory integration, the neocortex-like model reproduce a similar neuron-activation pattern which is obtained in the hippocampus-like model, and still keeping the previous activation pattern when processing previous SWR-like signals.

5 Conclusion

In this work, we propose a general and bio-inspired incremental learning framework named SWRL. The proposed SWRL simulates the collaborative mechanism between the

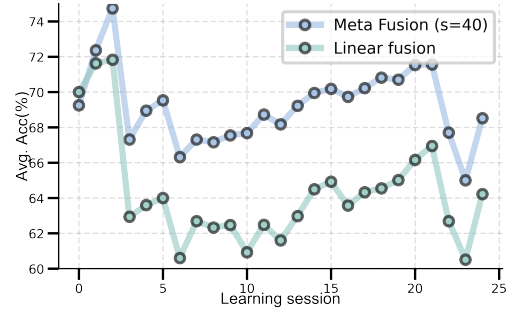


Figure 5: Comparison of incremental learning results using different memory integration settings. The notation "s = 40" indicates that 40 SWR-like signals are generated for each learning session.

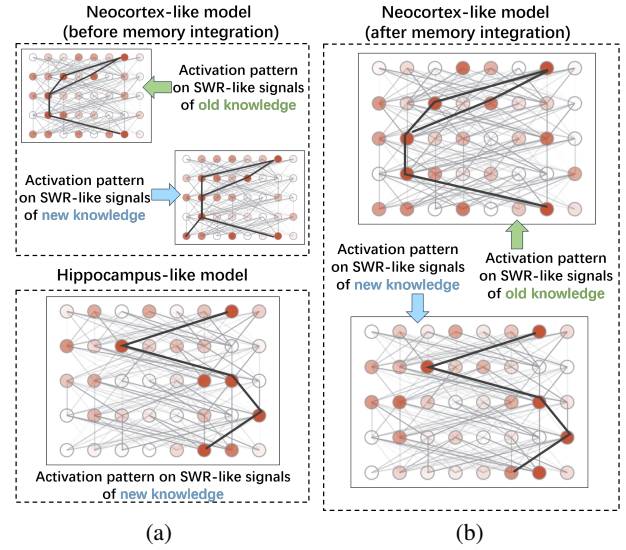


Figure 6: Neuron-circuits activation results before and after memory integration (CIFAR100-TIL-s10, Resnet-18).

hippocampus and neocortex by employing a two stage learning framework, including the fast learning stage and meta-weighted memory integration guided by SWR-like signals. SWRL is a task-agnostic, architecture-agnostic, and buffer-free methods. Comprehensive experiments demonstrate its superior performance in different incremental learning scenarios, across various deep neural network architectures.

Acknowledgments

This work is supported by the National Natural Science Foundation of China (T2521006 and 61973311).

Contribution Statement

Yi Sun and Xiaochang Hu contributed equally to this work. Jian Li and Jiting Li are corresponding authors.

References

- [Apolinario *et al.*, 2025] Marco P. E. Apolinario, Sakshi Choudhary, and Kaushik Roy. Code-cl: Conceptor-based gradient projection for deep continual learning. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 775–784, October 2025.
- [Bellitto *et al.*, 2024] Giovanni Bellitto, Federica Proietto Salanitri, Matteo Pennisi, Matteo Boschini, Lorenzo Bonicelli, Angelo Porrello, Simone Calderara, Simone Palazzo, and Concetto Spampinato. Saliency-driven experience replay for continual learning. In *Advances in Neural Information Processing Systems*, volume 37, 2024.
- [Bian *et al.*, 2024] Ang Bian, Wei Li, Hangjie Yuan, Chengrong Yu, Mang Wang, Zixiang Zhao, Aojun Lu, Pengliang Ji, and Tao Feng. Make continual learning stronger via c-flat. In *Advances in Neural Information Processing Systems*, volume 37, 2024.
- [Chen *et al.*, 2024] Jinpeng Chen, Runmin Cong, Yuxuan Luo, Horace Ho Shing Ip, and Sam Kwong. Strike a balance in continual panoptic segmentation. In *European Conference on Computer Vision*, pages 126–142. Springer, 2024.
- [Dosovitskiy *et al.*, 2021] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. In *The Ninth International Conference on Learning Representations*, 2021.
- [Feng *et al.*, 2023] Haozhe Feng, Zhaorui Yang, Hesun Chen, Tianyu Pang, Chao Du, Minfeng Zhu, Wei Chen, and Shuicheng Yan. Cosda: Continual source-free domain adaptation. *arXiv preprint arXiv:2304.06627*, 2023.
- [Gao *et al.*, 2024a] Chen Gao, Xiaochong Lan, Nian Li, Yuan Yuan, Jingtao Ding, Zhilun Zhou, Fengli Xu, and Yong Li. Large language models empowered agent-based modeling and simulation: A survey and perspectives. *Humanities and Social Sciences Communications*, 11(1):1–24, 2024.
- [Gao *et al.*, 2024b] Zhanxin Gao, Jun Cen, and Xiaobin Chang. Consistent prompting for rehearsal-free continual learning. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 28463–28473, 2024.
- [He *et al.*, 2016] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.
- [Kang *et al.*, 2025] Borui Kang, Lei Wang, Zhiping Wu, Tao Feng, Yawen Li, Yang Gao, and Wenbin Li. Dynamic multi-layer null space projection for vision-language continual learning. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 2077–2086, 2025.
- [Kao *et al.*, 2021] Ta-Chu Kao, Kristopher Jensen, Gido van de Ven, Alberto Bernacchia, and Guillaume Hennequin. Natural continual learning: success is a journey, not (just) a destination. *Advances in neural information processing systems*, 34:28067–28079, 2021.
- [Ke *et al.*, 2025] Hualong Ke, Jiangming Shi, Yachao Zhang, Fangyong Wang, Yuan Xie, and Yanyun Qu. Task-aware prompt gradient projection for parameter-efficient tuning federated class-incremental learning. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 2631–2641, October 2025.
- [Kim *et al.*, 2024] Junsu Kim, Hoseong Cho, Jihyeon Kim, Yihalem Yimolal Tiruneh, and Seungryul Baek. Sd-dgr: Stable diffusion-based deep generative replay for class incremental object detection. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 28772–28781, June 2024.
- [Kirkpatrick *et al.*, 2017] James Kirkpatrick, Razvan Pascanu, Neil Rabinowitz, Joel Veness, Guillaume Desjardins, Andrei A Rusu, Kieran Milan, John Quan, Tiago Ramalho, Agnieszka Grabska-Barwinska, et al. Overcoming catastrophic forgetting in neural networks. *Proceedings of the national academy of sciences*, 114(13):3521–3526, 2017.
- [Krizhevsky *et al.*, 2009] Alex Krizhevsky, Geoffrey Hinton, et al. Learning multiple layers of features from tiny images. 2009.
- [Krukowski *et al.*, 2025] Patryk Krukowski, Anna Bielawska, Kamil Książek, Paweł Wawrzyński, Paweł Batorski, and Przemysław Spurek. Hint: Hypernetwork approach to training weight interval regions in continual learning. *Information Sciences*, 717:122261, 2025.
- [Kudithipudi *et al.*, 2022] Dhireesha Kudithipudi, Mario Aguilar-Simon, Jonathan Babb, Maxim Bazhenov, Douglas Blackiston, Josh Bongard, Andrew P Brna, Suraj Chakravarthi Raja, Nick Cheney, Jeff Clune, et al. Biological underpinnings for lifelong learning machines. *Nature Machine Intelligence*, 4(3):196–210, 2022.
- [Li *et al.*, 2025] Mei Li, Yuxiang Lu, Qinyan Dai, Suizhi Huang, Yue Ding, and Hongtao Lu. Became: Bayesian continual learning with adaptive model merging. *arXiv preprint arXiv:2504.02666*, 2025.
- [Liang and Li, 2024] Yan-Shuo Liang and Wu-Jun Li. In-flora: Interference-free low-rank adaptation for continual learning. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 23638–23647, 2024.
- [Lin *et al.*, 2022] Guoliang Lin, Hanlu Chu, and Hanjiang Lai. Towards better plasticity-stability trade-off in incremental learning: A simple linear connector. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 89–98, 2022.
- [Lomonaco and Maltoni, 2017] Vincenzo Lomonaco and Davide Maltoni. Core50: a new dataset and benchmark for continuous object recognition. In Sergey Levine, Vincent

- Vanhoecke, and Ken Goldberg, editors, *Proceedings of the 1st Annual Conference on Robot Learning*, volume 78 of *Proceedings of Machine Learning Research*, pages 17–26. PMLR, 13–15 Nov 2017.
- [Lu et al., 2024] Yue Lu, Shizhou Zhang, De Cheng, Yinghui Xing, Nannan Wang, Peng Wang, and Yanning Zhang. Visual prompt tuning in null space for continual learning. In *Advances in Neural Information Processing Systems*, volume 37, 2024.
- [Maracani et al., 2021] Andrea Maracani, Umberto Michieli, Marco Toldo, and Pietro Zanuttigh. Recall: Replay-based continual learning in semantic segmentation. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 7026–7035, 2021.
- [Masse et al., 2018] Nicolas Y Masse, Gregory D Grant, and David J Freedman. Alleviating catastrophic forgetting using context-dependent gating and synaptic stabilization. *Proceedings of the National Academy of Sciences*, 115(44):E10467–E10475, 2018.
- [Norman et al., 2019] Yitzhak Norman, Erin M Yeagle, Simon Khuvis, Michal Harel, Ashesh D Mehta, and Rafael Malach. Hippocampal sharp-wave ripples linked to visual episodic recollection in humans. *Science*, 365(6454):eaax1030, 2019.
- [Notin et al., 2024] Pascal Notin, Nathan Rollins, Yarin Gal, Chris Sander, and Debora Marks. Machine learning for functional protein design. *Nature biotechnology*, 42(2):216–228, 2024.
- [Qiao et al., 2025] Jingyang Qiao, Zhizhong Zhang, Xin Tan, Yanyun Qu, Wensheng Zhang, Zhi Han, and Yuan Xie. Gradient projection for continual parameter-efficient tuning. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2025.
- [Rebuffi et al., 2017] Sylvestre-Alvise Rebuffi, Alexander Kolesnikov, Georg Sperl, and Christoph H Lampert. icarl: Incremental classifier and representation learning. In *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition*, pages 2001–2010, 2017.
- [Shi and Wang, 2024] Haizhou Shi and Hao Wang. A unified approach to domain incremental learning with memory: Theory and algorithm. *Advances in Neural Information Processing Systems*, 36, 2024.
- [van de Ven et al., 2022] Guido M van de Ven, Tinne Tuytelaars, and Andreas S Tolias. Three types of incremental learning. *Nature Machine Intelligence*, 4(12):1185–1197, 2022.
- [Wang et al., 2021] Shipeng Wang, Xiaorong Li, Jian Sun, and Zongben Xu. Training networks in null space of feature covariance for continual learning. In *Proceedings of the IEEE/CVF conference on Computer Vision and Pattern Recognition*, pages 184–193, 2021.
- [Wang et al., 2022a] Yabin Wang, Zhiwu Huang, and Xiaopeng Hong. S-prompts learning with pre-trained transformers: An occam’s razor for domain incremental learning. *Advances in Neural Information Processing Systems*, 35:5682–5695, 2022.
- [Wang et al., 2022b] Zifeng Wang, Zizhao Zhang, Sayna Ebrahimi, Ruoxi Sun, Han Zhang, Chen-Yu Lee, Xiaoqi Ren, Guolong Su, Vincent Perot, Jennifer Dy, et al. Dualprompt: Complementary prompting for rehearsal-free continual learning. In *European Conference on Computer Vision*, pages 631–648. Springer, 2022.
- [Wang et al., 2022c] Zifeng Wang, Zizhao Zhang, Chen-Yu Lee, Han Zhang, Ruoxi Sun, Xiaoqi Ren, Guolong Su, Vincent Perot, Jennifer Dy, and Tomas Pfister. Learning to prompt for continual learning. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 139–149, 2022.
- [Wang et al., 2024] Liyuan Wang, Xingxing Zhang, Hang Su, and Jun Zhu. A comprehensive survey of continual learning: Theory, method and application. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2024.
- [Wang et al., 2025] Qiang Wang, Xiang Song, Yuhang He, Jizhou Han, Chenhao Ding, Xinyuan Gao, and Yihong Gong. Boosting domain incremental learning: Selecting the optimal parameters is all you need. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pages 4839–4849, 2025.
- [Wu et al., 2017] Jiayu Wu, Qixiang Zhang, and Guoxi Xu. Tiny imagenet challenge. *Technical report*, 2017.
- [Xie et al., 2025] Guanglei Xie, Yi Sun, Xin Xu, Hao Fu, Yifei Shi, and Xiaochang Hu. Ni-sscl: A neuroplasticity-inspired method for semi-supervised continual learning. *IEEE Transactions on Cognitive and Developmental Systems*, 2025.
- [Yang et al., 2024] Chengyi Yang, Mingda Dong, Xiaoyue Zhang, Jiayin Qi, and Aimin Zhou. Introducing common null space of gradients for gradient projection methods in continual learning. In *Proceedings of the 32nd ACM International Conference on Multimedia*, MM ’24, page 5489–5497, New York, NY, USA, 2024. Association for Computing Machinery.
- [Yu et al., 2024] Jiazuo Yu, Yunzhi Zhuge, Lu Zhang, Ping Hu, Dong Wang, Huchuan Lu, and You He. Boosting continual learning of vision-language models via mixture-of-experts adapters. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 23219–23230, June 2024.
- [Zhang et al., 2023] Xingxuan Zhang, Renzhe Xu, Han Yu, Hao Zou, and Peng Cui. Gradient norm aware minimization seeks first-order flatness and improves generalization. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 20247–20257, 2023.
- [Zhou et al., 2025] Da-Wei Zhou, Zi-Wen Cai, Han-Jia Ye, Lijun Zhang, and De-Chuan Zhan. Dual consolidation for pre-trained model-based domain-incremental learning. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pages 20547–20557, 2025.