

Dual Consolidation for Pre-Trained Model-Based Domain-Incremental Learning

Da-Wei Zhou, Zi-Wen Cai, Han-Jia Ye^(✉), Lijun Zhang, De-Chuan Zhan

School of Artificial Intelligence, Nanjing University

National Key Laboratory for Novel Software Technology, Nanjing University

{zhoudw, caizw, yehj, zhanglj, zhancd}@lamda.nju.edu.cn

Abstract

Domain-Incremental Learning (DIL) involves the progressive adaptation of a model to new concepts across different domains. While recent advances in pre-trained models provide a solid foundation for DIL, learning new concepts often results in the catastrophic forgetting of pre-trained knowledge. Specifically, sequential model updates can overwrite both the representation and the classifier with knowledge from the latest domain. Thus, it is crucial to develop a representation and corresponding classifier that accommodate all seen domains throughout the learning process. To this end, we propose DUAL Consolidation (DUCT) to unify and consolidate historical knowledge at both the representation and classifier levels. By merging the backbone of different stages, we create a representation space suitable for multiple domains incrementally. The merged representation serves as a balanced intermediary that captures task-specific features from all seen domains. Additionally, to address the mismatch between consolidated embeddings and the classifier, we introduce an extra classifier consolidation process. Leveraging class-wise semantic information, we estimate the classifier weights of old domains within the latest embedding space. By merging historical and estimated classifiers, we align them with the consolidated embedding space, facilitating incremental classification. Extensive experimental results on four benchmark datasets demonstrate DUCT's state-of-the-art performance. Code is available at <https://github.com/Estrella-fugaz/CVPR25-Duct>.

1. Introduction

Recent years have seen the rise of deep learning, demonstrating its strong potential in real-world applications [15, 17, 21, 42, 43, 76]. However, in a dynamic and ever-changing world, data often evolves in a stream format [1, 30], necessitating continuous updates with data from new

domains [19]. For instance, autonomous vehicles must handle driving tasks across different seasons and weather conditions [10], and face recognition systems are supposed to recognize users despite the varying lighting conditions [81]. Correspondingly, Domain-Incremental Learning (DIL) [60] addresses this challenge by absorbing new knowledge from new data distributions while preserving existing knowledge. However, sequentially updating the model with new domains often biases the embedding [77] and classifier [73, 80] modules toward the latest domain, leading to catastrophic forgetting [14, 40] of previous knowledge. Therefore, developing methods for learning new domains without forgetting existing knowledge becomes a critical challenge for the machine learning community.

To combat catastrophic forgetting, recent works leverage strong pre-trained models (PTMs) [20] as the initialization for DIL [67–69]. PTMs' robust feature representation abilities provide a powerful starting point, showing promising results in DIL benchmarks. These approaches freeze the pre-trained backbone to preserve existing knowledge and append lightweight modules [13, 27, 32] to capture new patterns. However, since incoming domains will also overwrite the appended modules, the representations still suffer from forgetting, leading to corresponding bias in the classifier.

In DIL, there are two primary sources of forgetting: the overwriting of *features* and *classifiers*. Continually adjusting the features to capture the latest task biases the representations toward the newest domain. Assuming the previous task involves real-world photos while the subsequent one consists of quick-draw images, the features will be overwhelmingly adjusted to highlight non-objective details. This may hinder the model from distinguishing the previous domain, resulting in forgetting. In light of this, an ideal feature space in DIL should fit all domains and equally highlight the domain-specific features for all seen tasks. However, as features become less generalizable due to overwriting, the classifier also becomes biased toward the latest domain. Since the classifier plays the pivotal role of mapping embeddings to classes, using the biased classifier can cause detrimental effects on decision-making. Consequently, it

[†]Correspondence to: Han-Jia Ye (yehj@lamda.nju.edu.cn)

is imperative to consolidate existing knowledge at both the representation and classifier levels to resist forgetting.

There are two main challenges to achieving this goal.

1) Building a unified embedding space that suits all tasks continually. Since the training data arrives in a stream format, we cannot access all training instances simultaneously, preventing us from achieving the ‘oracle embedding’ that favors all seen domains. **2)** Calibrating the prediction of biased classifiers. Even if we achieve a holistic embedding suitable for all tasks, we still cannot build an accurate mapping between features and classes due to the lack of previous instances. Consequently, a classifier consolidation process is needed to align the classifier weights with the changing features continually.

In this paper, we propose DUAL Consolidation (DUCT) to address these challenges. To counter feature drift, we introduce a representation merging technique that continuously integrates historical backbones. This merged embedding consolidates task-specific information from all previous domains, providing informative features without forgetting. Consequently, we obtain non-forgettable features incrementally and resist feature-level forgetting. To map updated features to classes, we design a classifier consolidation process that merges historical classifier weights with calibrated weights. By extracting class-wise semantic relationships in the joint space, we develop a classifier transport mechanism that estimates classifiers of previous domain classes in the latest embedding space. Through consolidating features and classifiers in this coordinated way, DUCT balances all seen domains and robustly resists forgetting, achieving strong performance on various benchmarks.

2. Related Work

Continual Learning/Incremental Learning aims to incorporate new knowledge within streaming data [14, 37, 66]. It can be further classified into three subcategories, *i.e.*, task-incremental learning (TIL) [14], class-incremental learning (CIL) [37, 86], and domain-incremental learning (DIL) [60]. Among them, TIL and CIL are faced with emerging new classes and are required to learn them without forgetting. By contrast, DIL mainly focuses on the scenario where label space is fixed while incoming data involves new domains [41, 55, 67]. In continual learning, typical solutions include using data replay [7, 8, 26, 33, 34, 50, 52, 82] to review former knowledge, using knowledge distillation [22, 31, 51] or parameter regularization [5, 6, 12, 24, 28, 78] to maintain existing knowledge. Recent advances also resort to bias correction [2, 9, 11, 23, 73, 80] to rectify the model bias or expand the network structure as data evolves [4, 54, 64, 65, 74, 83, 84].

Domain-Incremental Learning with PTMs: With the rapid development of pre-training techniques, PTMs provide a strong initialization for DIL models to boost the fea-

ture generalizability [36, 48, 57, 59, 67–69, 75]. Current PTM-based DIL methods mainly resort to visual prompt tuning [27] to adjust the pre-trained features. With the pre-trained features frozen, it learns a prompt pool as external knowledge and selects instance-specific prompts to encode task information. L2P [69] utilizes the query-key matching mechanism to select instance-specific prompts, while DualPrompt [68] extends it by learning task-specific and instance-specific prompts jointly. By contrast, S-Prompts [67] learns domain-specific prompts for DIL and retrieves prompts via KNN search. To alleviate the prompt selection cost, CODA-Prompt [57] reformulates this selection step as an attention-based weighted combination process. There are also some methods that directly build classifiers upon the pre-trained features [39, 87].

Model Merging: Recent advances in pre-training have made weight interpolation among different models a useful technique in model editing [3, 18, 25, 38, 71, 72]. Most work focuses on simple averaging of multiple models, aiming to enhance model in terms of its performance on the single task and robustness to out-of-distribution data. This paper considers merging pre-trained weights and task vectors to cultivate multitask representation for all seen domains.

3. From Old Domains to New Domains

3.1. Domain-Incremental Learning

In domain-incremental learning, the model is faced with the data stream containing a sequence of tasks $\{\mathcal{D}^1, \mathcal{D}^2, \dots, \mathcal{D}^B\}$, where $\mathcal{D}^b = \{\mathcal{X}_b, \mathcal{Y}_b\}$ is the dataset of the b -th training task, *i.e.*, $\mathcal{X}_b = \{\mathbf{x}_i\}_{i=1}^{n_b}$ and $\mathcal{Y}_b = \{y_i\}_{i=1}^{n_b}$. Each training instance $\mathbf{x}_i \in \mathbb{R}^D$ belongs to class $y_i \in Y$. This paper follows the **exemplar-free** setting [69, 88], *i.e.*, during the b -th training stage, we can only access data in the current dataset $\mathcal{D}^b$. In DIL, the label space Y is unchanged throughout the learning process, while the distribution of input instance keeps changing from domain to domain, *i.e.*, $p(\mathcal{X}_b) \neq p(\mathcal{X}_{b'})$ for $b \neq b'$. The target of DIL is to fit a model $f(\mathbf{x})$ that can discriminate the classes among any seen domains:

$$f^* = \operatorname{argmin}_{f \in \mathcal{H}} \mathbb{E}_{(\mathbf{x}, y) \sim \mathcal{D}_1^1 \cup \dots \mathcal{D}_t^1} \mathbb{I}(y \neq f(\mathbf{x})), \quad (1)$$

where $\mathcal{H}$ is the hypothesis space, $\mathbb{I}(\cdot)$ is the indicator function which outputs 1 if the expression holds and 0 otherwise. $\mathcal{D}_t^b$ denotes the data distribution of task b . Consequently, DIL models are supposed to classify instances of new domains while not forgetting previous ones.

Following [57, 67–69], we assume a pre-trained Vision Transformer (ViT) [16] is available as the initialization for $f(\mathbf{x})$. For simplicity, we decouple the network structure into the embedding function $\phi(\cdot) : \mathbb{R}^D \rightarrow \mathbb{R}^d$ (*i.e.*, the final [CLS] token) and the linear classifier $W \in \mathbb{R}^{d \times |Y|}$. In

this way, the output is denoted as $f(\mathbf{x}) = W^\top \phi(\mathbf{x})$, and we utilize a cosine classifier in this paper. The classifier is further decoupled into $W = [\mathbf{w}_1, \mathbf{w}_2, \dots, \mathbf{w}_{|Y|}]$, where $\mathbf{w}_j$ denotes the classifier weight of the j -th class. Although the label space is static during the learning process, data distribution of the same class across different domains can yield significant domain gaps. Hence, we follow existing works [57, 67–69] to scale the classifier up to $W \in \mathbb{R}^{d \times b|Y|}$ when the b -th task emerges, and the final prediction is made by $(\arg \max_i \mathbf{w}_i^\top \phi(\mathbf{x})) \bmod |Y|$.

3.2. Baselines in Domain-Incremental Learning

In domain-incremental learning, a naive solution to handle the incoming datasets of new domains is to directly optimize the embedding and classifier:

$$\min_{W \cup \phi} \sum_{(\mathbf{x}, y) \in D^b} \ell(W^\top \phi(\mathbf{x}), y). \quad (2)$$

With pre-trained weights as initialization, Eq. (2) can quickly capture the discriminative features within the new domain. However, since the embedding $\phi(\cdot)$ keeps changing across different domains, it quickly loses the generalization ability to previous domains and suffers from forgetting.

To resist feature-level forgetting, a feasible solution is to freeze the pre-trained weights and append lightweight modules to encode domain-specific knowledge, e.g., prompt [27]. A representative work L2P [69] organizes a set of prompts as the prompt pool ($\mathcal{P}$) and optimizes:

$$\min_{\mathcal{P} \cup W} \sum_{(\mathbf{x}, y) \in D^b} \ell(W^\top \bar{\phi}(\mathbf{x}; \mathcal{P}), y) + \mathcal{L}_{\mathcal{P}}, \quad (3)$$

where $\bar{\phi}(\mathbf{x}; \mathcal{P})$ denotes the prompted feature representation with backbone frozen, and $\mathcal{L}_{\mathcal{P}}$ denotes the prompt selection loss [69]. Eq. (3) dynamically retrieves instance-specific prompts to adjust the representations and learns the classifier to map the features to corresponding classes. Since the backbone weights are frozen, it can alleviate the representation drift during updating.

Discussions: Although Eq. (2) and Eq. (3) adopt different policies for model updating, both of them suffer from the *feature-level* forgetting. Specifically, tuning the model via Eq. (2) can encode task-specific knowledge into the embedding, but also comes with the risk of features being fully overwritten by new domains since there is no restriction on retaining previous knowledge. Besides, although Eq. (3) freezes the pre-trained weights, the trainable prompts appended may also incur forgetting. Since the prompted feature relies on instance-specific prompt selection, optimizing the prompts for new domains still has the risk of overwriting previous ones. With the features being biased toward the latest domain, the classifier is adjusted accordingly. Consequently, the classifiers of previous tasks are incompatible with the unstable and ever-changing features, resulting in the *classifier-level* forgetting.

4. DUCT: Dual Consolidation for Domain-Incremental Learning

Observing the risk of feature-level and classifier-level overwriting, we aim to resist forgetting in DIL from two aspects. Firstly, alleviating feature-level forgetting requires obtaining a unified embedding space at a low cost. Since sequentially updating the backbone or prompts will both result in forgetting, a proper solution is needed to make full use of all historical features and consolidate them into a unified one. Secondly, as the embedding space changes from task to task, the mismatch between representation and classifier becomes more severe as the data evolves. Consequently, a classifier consolidation step is needed to calibrate them to be compatible with the embeddings.

4.1. Representation Consolidation

Observing the challenge of striking a balance between all seen domains, we need to opt for another feature updating policy to avoid sequential overwriting and resist forgetting. Let us assume an ideal scenario where we can separately train each model for each incoming domain, obtaining a set of models $\{\phi_1(\cdot), W_1\}, \dots, \{\phi_B(\cdot), W_B\}$. Those models are obtained by optimizing the same PTM via Eq. (2), thus having the discriminability for each specific domain and can be seen as the ‘domain expert.’ In an oversimplified scenario where we know which domain the instance is from, we can utilize the corresponding expert for prediction. However, since such auxiliary information is unavailable in DIL, we need to obtain an *omnipotent* embedding space that suits all domains. In this way, there is no need to decide which embedding to use since the omnipotent embedding is strong enough to capture all domain-specific features.

Correspondingly, we take inspiration from model merging [25, 49] and cognitive-developmental theory [47] that an ideal model for multiple domains can be achieved by combining multiple *task vectors*. We denote the relative change of the embedding function as $\delta_{\phi_i} = \phi_i - \phi_0$, where ϕ_0 is the weight of the pre-trained model. Hence, δ_{ϕ_i} represents the relative change of the weights in the i -th domain, and we can build the unified embedding via:

$$\phi_i^m = \phi_0 + \alpha_\phi \sum_{k=1}^i \delta_{\phi_k}, \quad (4)$$

where α_ϕ, ϕ_i^m refers to the weight of task vector δ_{ϕ_k} and the merged backbone respectively. For task vector start from the same pre-trained weight, [25, 79] verify the task vectors share low similarity due to the semantic gap between different domains. Consequently, adding such weights enables the merged model to highlight all task-specific features, and Eq. (4) provides a feasible way to obtain the universal feature representation that suits all seen domains.

Representation consolidation with task similarity: Although Eq. (4) provides a simple way to consolidate multiple task vectors to build a unified embedding, it still lacks

consideration in measuring task-wise similarities. For example, if two distinct domains are more similar, highlighting those features would be more helpful for recognizing corresponding domains. Hence, we introduce task-similarity ($\mathbf{Sim}_{0,i}$) into the consolidation process.

To measure the similarity, a naive way is to directly measure the similarity between the embedding function of pre-trained model ϕ_0 and that of the subsequent model ϕ_i . However, since the weights of pre-trained model are with millions of dimensions, the similarity in such space is ineffective due to the curse of dimensionality. To this end, we utilize the current training task as an indicator of task similarity. Specifically, for every class in $\mathcal{D}^b$, we utilize the backbone ϕ_i to extract class centers in its embedding space:

$$\mathbf{c}_p^i = \sum_{j=1}^{|\mathcal{D}^b|} \mathbb{I}(y_j = p) \phi_i(\mathbf{x}_j) / \sum_{j=1}^{|\mathcal{D}^b|} \mathbb{I}(y_j = p). \quad (5)$$

Afterward, we can obtain two sets of class centers in the embedding space, *i.e.*, $\mathbf{C}^0 = [\mathbf{c}_1^0, \mathbf{c}_2^0, \dots, \mathbf{c}_{|Y|}^0]$, $\mathbf{C}^i = [\mathbf{c}_1^i, \mathbf{c}_2^i, \dots, \mathbf{c}_{|Y|}^i]$. Since those class centers are representative points in the embedding space, we calculate the pairwise class similarity to indicate task similarities:

$$\mathbf{Sim}_{0,i} = \frac{1}{|Y|} \sum_{j=1}^{|Y|} \text{sim}(\mathbf{c}_j^0, \mathbf{c}_j^i), \quad (6)$$

where we utilize cosine similarity to calculate center-wise similarity $\text{sim}(\cdot, \cdot)$. We then introduce task similarity into Eq. (4), and the unified embedding is denoted as:

$$\phi_i^m = \phi_0 + \alpha_\phi \sum_{k=1}^i \mathbf{Sim}_{0,k} \delta_{\phi_k}, \quad (7)$$

Effect of representation consolidation: Figure 1 (top) visualizes the representation consolidation process, where the merged backbone can unify multiple domains. Through Eq. (7), we are able to build the joint embedding space that suits all tasks by combining them in the parameter space. Since tasks in DIL emerge one-by-one, the representation consolidation process can also be made *incrementally*, *i.e.*, $\phi_i^m = \phi_{i-1}^m + \alpha_\phi \mathbf{Sim}_{0,i} \delta_{\phi_i}$. In this way, we can get rid of the extensive memory cost and only keep at most two backbones in memory. In each training stage, we first fine-tune the model via Eq. (2) and then conduct representation consolidation to aggregate the representations. In this way, we obtain a unified embedding space across all seen tasks.

4.2. Classifier Consolidation

In Eq. (7), we design the representation consolidation process, which aggregates multiple fine-tuned models by accumulating task vectors. However, a fatal problem still exists, *i.e.*, there is a *mismatch* between the classifiers and the consolidated embeddings. Since the classifiers are optimized to align the embeddings with the corresponding class, the matching degree can decrease drastically when the embedding space is totally updated. Recalling the background information in Section 3.1 that we utilize an independent classifier for each domain, we need to design an extra classifier

consolidation process to align the classifiers with the updated embedding space. In each training stage, we denote the classifier for previous domains as ‘old classifier’ (W_o) and the classifier for the current domain as ‘new classifier’ (W_n). We then discuss how to align them with the consolidated features.

New Classifier Retraining: In each training stage, we have the training data $\mathcal{D}^b$ in hand. Hence, it is intuitive to coordinate the new classifier with the consolidated features via:

$$\min_{W_n} \sum_{(\mathbf{x}, y) \in \mathcal{D}^b} \ell(W_n^\top \bar{\phi}_i^m(\mathbf{x}), y), \quad (8)$$

where the consolidated feature $\bar{\phi}_i^m(\cdot)$ is frozen to resist further mismatch, and we can adjust the new classifier W_n to match the features using Eq. (8).

Old Classifier Transport: Though Eq. (8) aligns the new classifier to the latest embedding space, another issue remains, *i.e.*, the old classifier is also incompatible with the merged embedding space. Consequently, the previous knowledge of old domains shall be forgotten in the incremental learning process. If we have plenty of training instances of previous domains, a similar calibration process can be done as in Eq. (8). However, due to the exemplar-free restrictions, we cannot save any previous instances in the memory. Hence, we need to find another way to estimate the relative calibration weight for the old classifier. For example, if we have the retrained classifier to classify a ‘lion’ in the *clip art* style, it can be mostly reused to classify the lion in the *real photo* style. We call such a relationship ‘semantic information,’ and the goal is to utilize such information to assist old classifier alignment. We denote the calibrated classifier using semantic information as $\hat{W}_o = \mathcal{T}(W_n, \mathbf{S})$, *i.e.*, the estimated classifier is obtained by transforming the new classifier using semantic information $\mathbf{S}$. Hence, there remain two core problems to be solved: 1) How to define the transformation function $\mathcal{T}$? 2) How to define the semantic information $\mathbf{S}$?

Implementing $\mathcal{T}$ via Optimal Transport: The function $\mathcal{T}$ encodes the correlation between the set of features among two tasks. Considering that the weight matrix of a linear layer reveals the relative relationship among feature-class pairs and the final logit is the aggregation of all features, we can design the *recombination* of existing classifiers with a linear mapping $T \in \mathbb{R}^{\beta \times \gamma}$. T encodes the cross-task correlation between the current and the previous domains, where larger values in T denote higher class-wise similarity. Since the decision is made by matching the classifier with corresponding features, we can utilize and recombine the weights of similar classes in the current domain to obtain those of previous domains. For example, important features that discriminate lions in the clip art style should be assigned higher coefficients to help classify lions in the photo style and vice versa.

We denote $\mu_1 \in \Delta_\beta$ and $\mu_2 \in \Delta_\gamma$, where $\Delta_d =$

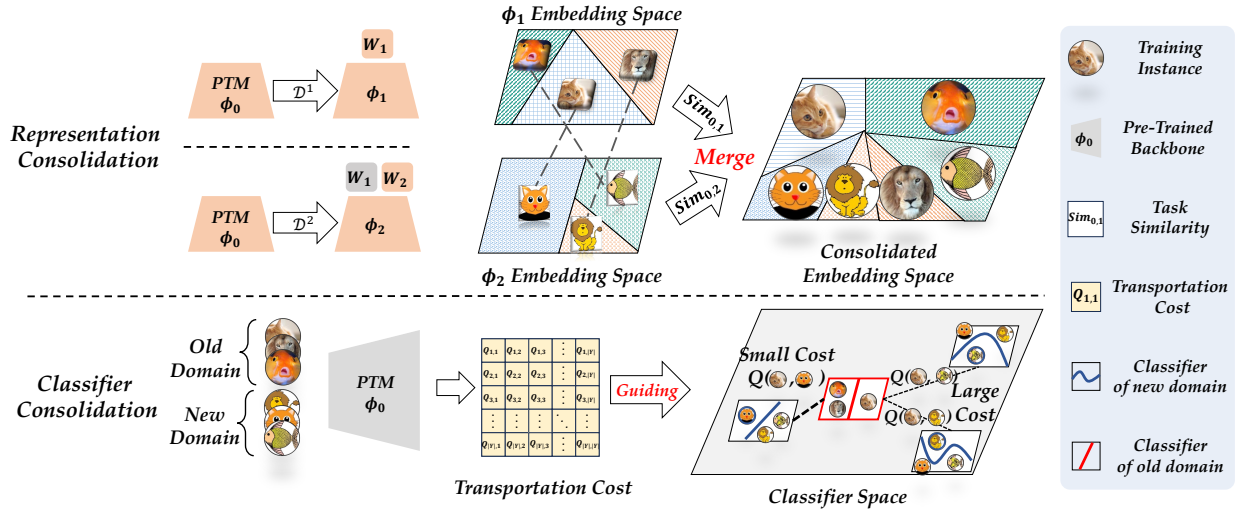


Figure 1. Illustration of DUCT. **Top:** Representation consolidation. We utilize the pre-trained model as initialization and optimize it for each domain, obtaining the task vectors. Afterward, we combine the pre-trained model and all seen task vectors to build the unified embedding space. **Bottom:** Classifier consolidation. To align the classifiers with consolidated features, we design the new classifier retraining and old classifier transport to consolidate classifiers. Class-wise semantic information is utilized in classifier transport.

$\{\mu : \mu \in \mathbb{R}_+^d, \mu^\top \mathbf{1} = 1\}$ is the d -dimensional simplex. μ_1 and μ_2 denote the importance of each class across domains, and we set them as uniform distributions. To map the distribution of β new classes to that of γ classes from old domains, a cost matrix $Q \in \mathbb{R}_+^{\beta \times \gamma}$ is further introduced to guide the transition. The larger weight of $Q_{i,j}$ indicates we need to pay more cost when reusing the classifier of i -th class to assist the j -th class. Consequently, the transport matrix T can be formulated as the coupling of two distributions, aiming to connect classes from different domains at the lowest transportation cost. Hence, T can be obtained via minimizing:

$$\min_T \langle T, Q \rangle \quad \text{s.t. } T\mathbf{1} = \mu_1, T^\top \mathbf{1} = \mu_2, T \geq 0, \quad (9)$$

which is the Kantorovich formulation of optimal transport (OT) [46, 63]. Optimizing Eq. (9) enables us to show how to move the probability mass across domains with low cost. Since the target is to reuse the retrained classifier W_n to adjust previous classifier W_o , we set $\mu_1 \in \Delta_{|Y|}$ and $\mu_2 \in \Delta_{|Y_o|}$ with uniform class marginals, where $|Y_o|$ denotes the number of classes in W_o . Solving Eq. (9) establishes the mapping relationship between different domains, and we can apply T to the new classifier to obtain the estimated old classifier, i.e., $\hat{W}_o = W_n T$.

Defining the Transportation Cost with Semantic Information: Solving Eq. (9) requires a proper definition of the cross-domain cost, i.e., Q . The higher cost indicates it is less effective to transport the classifier to the target class and vice versa. To this end, we design a simple yet effective way to measure such class-wise information. Specifically, before the training of each stage, we utilize the pre-trained backbone ϕ_0 to extract class centers in its embedding space via Eq. (5), i.e., $[c_1^0, c_2^0, \dots, c_{|Y|}^0]$. These class centers indicate the most representative embedding of the correspond-

ing classes. Since ϕ_0 is optimized with extensive datasets, we rely on its generalizability to reflect task-wise similarity. If two classes are similar, their embeddings should also be situated near each other. Consequently, we calculate the Euclidean distance between class centers as the transportation cost, i.e., $Q_{i,j} = \|c_i^0 - c_j^0\|_2^2$. Here, classes i and j are from different domains. With the pair-wise transportation cost, we are able to optimize Eq. (9) for the transportation plan T . Afterward, we utilize the interpolation between old classifiers and the transported classifier for consolidation:

$$W_o^m = (1 - \alpha_W)W_o + \alpha_W \hat{W}_o = (1 - \alpha_W)W_o + \alpha_W W_n T. \quad (10)$$

Effect of classifier consolidation: We visualize the classifier consolidation process in Figure 1 (bottom). The classifier consolidation process takes two steps, i.e., new classifier retraining and old classifier transport. The first step is designed to align the new classifiers with the unified embedding, while the second step aims to recombine the old classifiers with the calibrated classifier. In this way, we obtain a unified classifier compatible with the consolidated features, thus favoring the recognition in the unified embedding space.

4.3. Summary of DUCT

In DUCT, facing a new training task, we extract the class centers with ϕ_0 . Afterward, we optimize the model and separately consolidate the features and classifiers. The consolidated model $([W_o^m; W_n]^\top \phi_i^m)$ is then utilized for testing. We utilize Sinkhorn's algorithm [56] to solve the OT problem in Eq. (9). We summarize the training steps with pseudo code in the supplementary material.

Discussions on memory cost: During the training process, we need to keep two models in the memory, i.e., the histor-

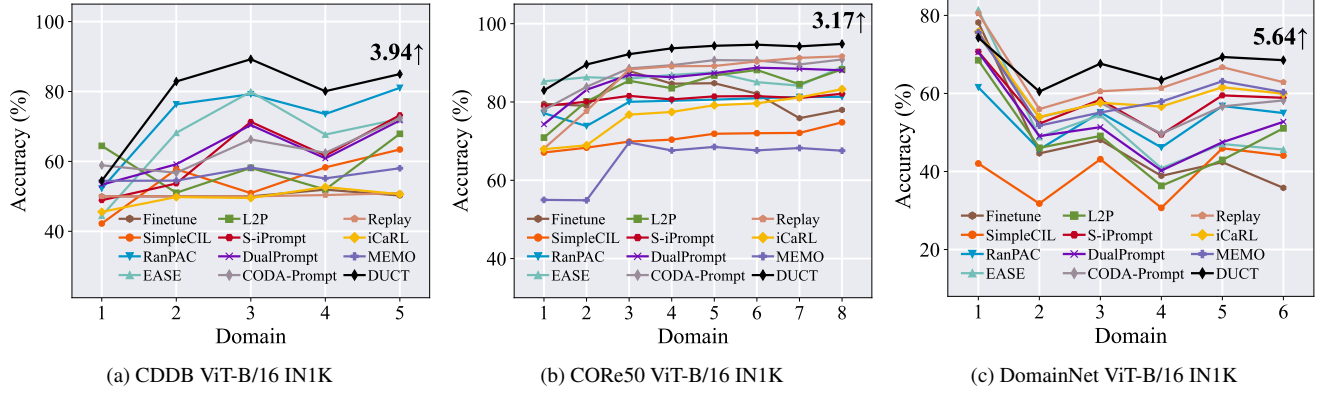


Figure 2. Incremental performance of different methods with the same pre-trained model. We report the performance gap after the last incremental stage between DUCT and the runner-up method at the end of the line.

Table 1. Average and last performance of different methods among five task orders. The best performance is shown in bold. All methods are implemented with ViT-B/16 IN1K. Methods with † indicate implemented with exemplars (10 per class).

Method	Office-Home		DomainNet		CORE50		CDDB	
	$\bar{\mathcal{A}}$	$\mathcal{A}_B$	$\bar{\mathcal{A}}$	$\mathcal{A}_B$	$\bar{\mathcal{A}}$	$\mathcal{A}_B$	$\bar{\mathcal{A}}$	$\mathcal{A}_B$
Finetune	78.32 $\pm$ 3.28	76.16 $\pm$ 1.39	28.17 $\pm$ 6.47	38.82 $\pm$ 7.65	75.44 $\pm$ 1.68	76.19 $\pm$ 2.36	52.08 $\pm$ 1.35	50.11 $\pm$ 1.62
Replay [†] [50]	84.23 $\pm$ 2.31	83.75 $\pm$ 0.68	64.78 $\pm$ 2.98	61.16 $\pm$ 1.19	85.56 $\pm$ 0.38	92.21 $\pm$ 0.63	66.91 $\pm$ 18.0	63.21 $\pm$ 11.6
iCaRL [†] [51]	81.66 $\pm$ 2.43	81.11 $\pm$ 1.23	59.89 $\pm$ 2.86	57.46 $\pm$ 2.31	74.43 $\pm$ 3.18	79.86 $\pm$ 2.96	68.43 $\pm$ 18.7	70.50 $\pm$ 16.5
MEMO [†] [84]	71.18 $\pm$ 2.76	63.09 $\pm$ 1.80	61.92 $\pm$ 5.39	58.41 $\pm$ 3.20	64.80 $\pm$ 3.16	68.24 $\pm$ 2.41	60.87 $\pm$ 13.4	58.09 $\pm$ 11.7
SimpleCIL [87]	75.69 $\pm$ 5.03	75.72 $\pm$ 0.00	42.95 $\pm$ 4.84	44.08 $\pm$ 0.00	70.92 $\pm$ 0.74	74.80 $\pm$ 0.00	60.80 $\pm$ 4.49	63.40 $\pm$ 0.00
L2P [69]	79.72 $\pm$ 4.19	80.03 $\pm$ 1.29	50.45 $\pm$ 4.10	48.72 $\pm$ 2.83	83.57 $\pm$ 0.35	87.87 $\pm$ 0.51	67.33 $\pm$ 5.98	64.45 $\pm$ 5.83
DualPrompt [68]	80.20 $\pm$ 3.81	80.85 $\pm$ 0.14	52.28 $\pm$ 3.35	50.46 $\pm$ 3.17	84.53 $\pm$ 0.89	87.27 $\pm$ 1.06	68.33 $\pm$ 7.52	71.41 $\pm$ 1.24
CODA-Prompt [57]	84.70 $\pm$ 2.94	85.07 $\pm$ 0.34	59.85 $\pm$ 4.49	59.99 $\pm$ 0.88	87.92 $\pm$ 0.41	91.57 $\pm$ 0.69	69.19 $\pm$ 6.11	74.18 $\pm$ 1.33
EASE [85]	81.16 $\pm$ 3.52	76.33 $\pm$ 2.16	50.50 $\pm$ 2.27	43.72 $\pm$ 1.70	86.30 $\pm$ 0.04	87.02 $\pm$ 1.21	67.78 $\pm$ 2.44	64.96 $\pm$ 8.36
RanPAC [39]	82.30 $\pm$ 3.34	82.28 $\pm$ 0.00	55.20 $\pm$ 3.93	54.80 $\pm$ 0.36	79.16 $\pm$ 0.55	81.38 $\pm$ 0.12	78.92 $\pm$ 4.85	80.48 $\pm$ 0.68
S-iPrompt [67]	81.50 $\pm$ 3.17	80.51 $\pm$ 0.21	61.16 $\pm$ 4.57	60.46 $\pm$ 0.90	81.95 $\pm$ 0.57	83.38 $\pm$ 0.70	68.51 $\pm$ 7.20	72.76 $\pm$ 0.43
DUCT	86.27$\pm$2.95	86.91$\pm$0.06	67.16$\pm$3.75	67.01$\pm$1.35	91.95$\pm$0.15	94.47$\pm$0.33	84.14$\pm$4.37	85.10$\pm$0.52

ical consolidated backbone ϕ_{i-1}^m and the online backbone ϕ_i for learning the current task. However, after the learning process of each domain, the current running model is merged into ϕ_i^m (Eq. (7)) without the need to be stored in memory. During inference, we utilize the consolidated backbone ϕ_i^m and classifiers $[W_o^m; W_n]$ for classification, which is the same as a single backbone. Consequently, DUCT can be fairly compared with others.

5. Experiment

In this section, we conduct experiments on benchmark datasets to compare DUCT to existing state-of-the-art methods. We also provide ablation studies on the components and parameter robustness to analyze the effect of different modules. Visualizations of the embedding space also verify the effectiveness of our proposed method.

5.1. Experimental Setup

Dataset: Following the benchmark settings [57, 67–69] in PTM-based domain-incremental learning, we evaluate the performance on **Office-Home** [62], **DomainNet** [45],

CORE50 [35], and **CDDB-Hard** [29]. Specifically, Office-Home has four domains, DomainNet has six domains, CORE50 has 11 domains, and CDDB-Hard has five domains. We consider five task orders to shuffle these domains in the DIL setting for a holistic evaluation and report the details in the supplementary material.

Comparison methods: In the comparison, we consider two types of methods, including exemplar-based methods, *i.e.*, Replay [50], iCaRL [51], MEMO [84], and SOTA exemplar-free DIL methods, *i.e.*, SimpleCIL [87], L2P [69], DualPrompt [68], CODA-Prompt [57], EASE [85], RanPAC [39], and S-iPrompt [67]. We utilize the **same pre-trained backbone for all compared methods**.

Implementation details: We deploy the experiments using PyTorch [44] and Pilot [58] on NVIDIA 4090. Following [69, 85], we consider two typical pre-trained weights, *i.e.*, ViT-B/16-IN21K and ViT-B/16-IN1K. Both are pre-trained with ImageNet21K [53], while the latter is further finetuned on ImageNet1K. We optimize DUCT using SGD optimizer with a batch size of 128 for 15 epochs. The learning rate is set to 0.001. We select 10 exemplars per

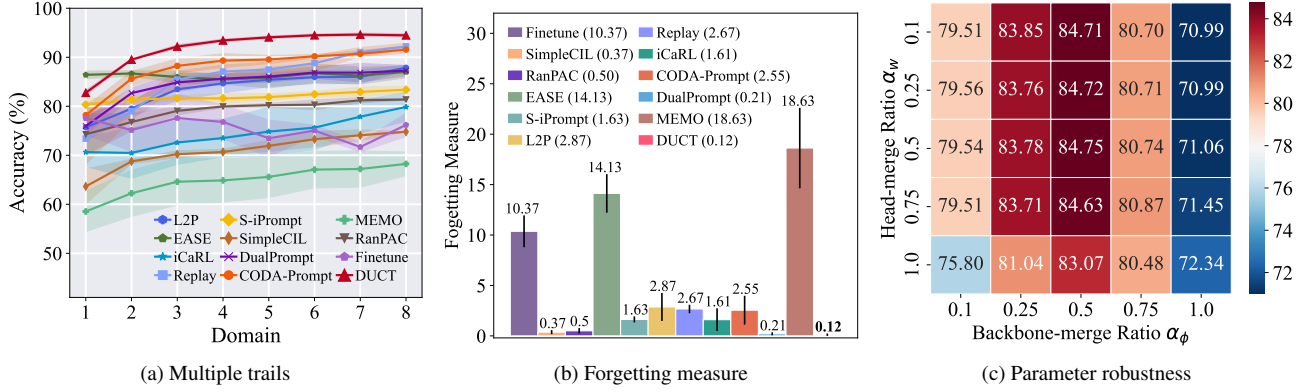


Figure 3. Further analysis on multiple task orders, forgetting measure, and parameter robustness. **(a):** Incremental performance of different methods on CORE50 with five task orders. The shadow indicates standard deviation. **(b):** Forgetting measure (**lower is better**) of different methods on CDDb dataset among five task orders. DUCT shows the least forgetting among all methods. **(c):** Average incremental performance with change of the consolidation ratios.

class for exemplar-based methods using herding [70] algorithm. In DUCT, we set the consolidation parameter $\alpha_\phi = 0.5, \alpha_w = 0.5$. The code will be made publicly available upon acceptance.

Performance Measure: Following [67, 69], we denote the accuracy among all seen domains after learning the b -th task as $\mathcal{A}_b$, we mainly consider $\mathcal{A}_B$ (the performance after learning the last stage) and $\bar{\mathcal{A}} = \frac{1}{B} \sum_{b=1}^B \mathcal{A}_b$ (the average performance among all incremental stages) for comparison. We also consider the forgetting measure [12] to measure the relative forgetting degree in DIL.

5.2. Benchmark Comparison

We report the results on benchmark datasets in Figure 2 and Table 1, with the remaining results provided in the supplementary material. As we can infer from these results, DUCT consistently outperforms other compared methods by 1 ~ 7% in the final accuracy. It must be noted that the differences of the starting point are due to the different tuning techniques, which cannot be directly aligned since the number of trainable parameters are different in those methods. Specifically, sequentially finetuning the model suffers from catastrophic forgetting, even starting with the pre-trained model. To compensate for the forgetting phenomena, we observe that several exemplar-based methods (Replay, iCaRL, and MEMO) show stable improvements to the baseline. However, since some of these works are designed for the class-incremental learning scenario, the expansion or distillation target may not be optimal for the current setting. We then compare DUCT to the methods specially designed for PTMs and find prompt-based methods (L2P, DualPrompt, CODA-Prompt, and S-iPrompt) yield inferior performance to ours. We also observe that DUCT is compatible with various pre-trained weights and shows stable improvements with ViT-B/16 IN1K and IN21K (see supplementary).

Besides, we also report the incremental performance among five task orders (average and standard deviation) in

Table 1 and Figure 3a. As we can infer from the table, DUCT works robustly across different task orders, showing stable improvements against other state-of-the-art methods.

5.3. Further Analysis

Forgetting Measure: Figure 2 and Table 1 mainly focus on the accuracy measure, which utilizes all seen domains to measure the relative performance. Apart from the accuracy measure, we also follow [67, 69] to utilize the forgetting measure, which captures the stability of existing knowledge among previous domains. As shown in Figure 3b, we report the forgetting measure among all compared methods on the CDDb dataset. As we can infer from the figure, several methods suffer from severe forgetting, *e.g.*, Finetune and MEMO, indicating that they are unsuitable for the domain-incremental learning scenario. We also observe that prompt-based methods freeze the backbone to resist feature drift, which tackles the forgetting phenomena. However, DUCT still shows the least forgetting (*i.e.*, 0.12) among all compared methods, indicating its strong performance.

Parameter Robustness: There are two major consolidation steps in DUCT, *i.e.*, the representation consolidation in Eq. (7) and the classifier consolidation in Eq. (10). These consolidation steps involve the hyper-parameters for merging a set of models/classifiers, *i.e.*, the backbone merge ratio α_ϕ and the classifier merge ratio α_w . In this section, we conduct ablations on the choice of these parameters to investigate the robustness of DUCT to their changes. Specifically, we choose α_ϕ and α_w among $\{0.1, 0.25, 0.5, 0.75, 1.0\}$, resulting in 25 parameter combinations. We conduct experiments with these parameter combinations on the Office-Home dataset and present the average performance in Figure 3c. As we can infer from the figure, DUCT is generally robust to parameter changes. Besides, since the merge parameters are designed to trade-off between old and new knowledge, we find either a small value (*e.g.*, 0.1) or a large value (*e.g.*, 1.0) works poorly. The poor performance of these marginal parameters also

Variations	$\bar{A}$	A_B
Baseline	66.56	63.40
Variation 1	80.36	74.17
Variation 2	81.41	76.82
Variation 3	85.42	80.31
DUCT	87.74	82.35

Table 2. Ablation study on different modules in DUCT.

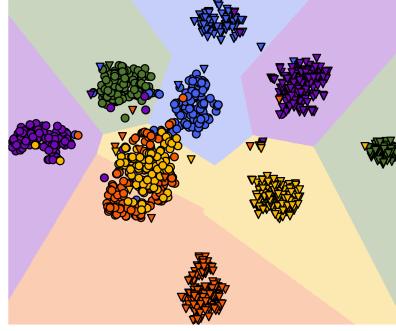


Figure 4. Before DUCT.

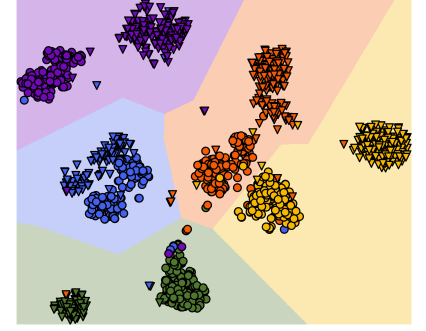


Figure 5. After DUCT.

Figure 6. Visualizations of the embedding space via t-SNE [61] on DomainNet. We visualize two incremental stages and utilize dots to represent the first domain and triangles for the second domain. The shadow region denotes the decision boundaries.

indicates the importance of the merging process since a small value or large value equals ignoring part of the important information during consolidation. By contrast, both of them prefer a mild value, *i.e.*, choosing them around $\{0.25, 0.5\}$ shows better performance. Hence, we suggest setting $\alpha_\phi = 0.5, \alpha_w = 0.5$ as the default setting.

Compositional Components: In this section, we conduct experiments to investigate the importance of each module in DUCT on CDDB dataset. As shown in Table 2, ‘Baseline’ denotes directly freezing the embedding and extracting class centers as the classifier, which shows poor performance due to the domain gap to the downstream domains. We then utilize Eq. (4) to sequentially merge the backbone to consolidate the representations and denote it as **Variation 1**. It shows that adding the representation consolidation process drastically improves the performance, indicating the superiority of such unified representation in DIL. However, when switching Eq. (4) to Eq. (7), we find **Variation 2** further improves the performance, indicating that task similarity is helpful in obtaining a universal embedding space. Afterward, we combine it with the classifier retraining process in Eq. (8), which is denoted as **Variation 3**. It shows that the mismatch between consolidated features and the classifiers weakens the performance, and aligning the new classifiers to the embedding helps DIL. When comparing DUCT to Variation 3, we find the old classifier transport is also vital to recover former knowledge and resist forgetting, which improves the final accuracy by 2%. Ablation studies verify the efficacy of different modules in DUCT.

Visualizations: In this section, we visualize the embedding space to show the effectiveness of DUCT on the DomainNet dataset. We consider a two-stage domain-incremental learning scenario, each containing five classes, and utilize t-SNE [61] to visualize the representation space before and after DUCT. We use dots to represent the classes in the first domain and triangles for classes in the second domain. We utilize the shadow regions to denote the decision boundary obtained by the classifier weights to represent different classes. As shown in Figure 4, although the embed-

ding space before DUCT shows competitive performance, there are two major flaws: 1) There exists the *confusion region* between yellow and red dots, indicating that previous classes are partially forgotten as data evolves. 2) The embedding space of the same class is still located in different regions, *e.g.*, the purple and green areas, which is not ideal for domain-incremental learning. However, when applying DUCT to consolidate the features and classifiers, the above problems are addressed in Figure 5. Specifically, since the consolidated embedding suits all tasks, the forgetting of previous domains is alleviated. Besides, the unified embedding space adequately clusters the same class of different domains together, favoring the final inference.

6. Conclusion

Domain-incremental learning is a desired ability for real-world learning systems to obtain new knowledge. In this paper, we propose dual consolidation (DUCT) to achieve this goal. Since the forgetting in DIL occurs in two aspects, *i.e.*, the representation and the classifier, we separately design the consolidation technique to handle them. Specifically, we consider continually merging the pre-trained model with domain-specific task vectors to achieve the unified embedding. To compensate for the mismatch between consolidated features and classifiers, we design the classifier consolidation process by introducing semantic-guided transport. Extensive experiments validate the effectiveness of DUCT.

Limitations: Possible limitations include the utilization of PTMs since representation consolidation relies on generalizable initialization. Future work includes extending DUCT to non-PTM scenarios.

Acknowledgments. This work is partially supported by NSFC (U23A20382, 62476123, 62376118, 62006112, 62250069), Fundamental Research Funds for the Central Universities (2024300373, 14380021), CCF-Tencent Rhino-Bird Open Research Fund RAGR20240101, Collaborative Innovation Center of Novel Software Technology and Industrialization.

References

- [1] Charu C Aggarwal. A survey of stream clustering algorithms. In *Data Clustering*, pages 231–258. Chapman and Hall/CRC, 2018. 1
- [2] Hongjoon Ahn, Jihwan Kwak, Subin Lim, Hyeonsu Bang, Hyojun Kim, and Taesup Moon. Ss-il: Separated softmax for incremental learning. In *ICCV*, pages 844–853, 2021. 2
- [3] Samuel K Ainsworth, Jonathan Hayase, and Siddhartha Srinivasa. Git re-basin: Merging models modulo permutation symmetries. In *ICLR*, 2023. 2
- [4] Rahaf Aljundi, Punarjay Chakravarty, and Tinne Tuytelaars. Expert gate: Lifelong learning with a network of experts. In *CVPR*, pages 3366–3375, 2017. 2
- [5] Rahaf Aljundi, Francesca Babiloni, Mohamed Elhoseiny, Marcus Rohrbach, and Tinne Tuytelaars. Memory aware synapses: Learning what (not) to forget. In *ECCV*, pages 139–154, 2018. 2
- [6] Rahaf Aljundi, Klaas Kelchtermans, and Tinne Tuytelaars. Task-free continual learning. In *CVPR*, pages 11254–11263, 2019. 2
- [7] Rahaf Aljundi, Min Lin, Baptiste Goujaud, and Yoshua Bengio. Gradient based sample selection for online continual learning. In *NeurIPS*, pages 11816–11825, 2019. 2
- [8] Jihwan Bang, Heesu Kim, YoungJoon Yoo, Jung-Woo Ha, and Jonghyun Choi. Rainbow memory: Continual learning with a memory of diverse samples. In *CVPR*, pages 8218–8227, 2021. 2
- [9] Eden Belouadah and Adrian Popescu. Il2m: Class incremental learning with dual memory. In *ICCV*, pages 583–592, 2019. 2
- [10] Mariusz Bojarski, Davide Del Testa, Daniel Dworakowski, Bernhard Firner, Beat Flepp, Prasoon Goyal, Lawrence D Jackel, Mathew Monfort, Urs Muller, Jiakai Zhang, et al. End to end learning for self-driving cars. *arXiv preprint arXiv:1604.07316*, 2016. 1
- [11] Francisco M Castro, Manuel J Marín-Jiménez, Nicolás Guil, Cordelia Schmid, and Karteek Alahari. End-to-end incremental learning. In *ECCV*, pages 233–248, 2018. 2
- [12] Arslan Chaudhry, Puneet K Dokania, Thalaiyasingam Ajanthan, and Philip HS Torr. Riemannian walk for incremental learning: Understanding forgetting and intransigence. In *ECCV*, pages 532–547, 2018. 2, 7
- [13] Shoufa Chen, Chongjian Ge, Zhan Tong, Jiangliu Wang, Yibing Song, Jue Wang, and Ping Luo. Adaptformer: Adapting vision transformers for scalable visual recognition. *NeurIPS*, 35:16664–16678, 2022. 1
- [14] Matthias De Lange, Rahaf Aljundi, Marc Masana, Sarah Parisot, Xu Jia, Aleš Leonardis, Gregory Slabaugh, and Tinne Tuytelaars. A continual learning survey: Defying forgetting in classification tasks. *TPAMI*, 44(7):3366–3385, 2021. 1, 2
- [15] Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Li Fei-Fei. Imagenet: A large-scale hierarchical image database. In *CVPR*, pages 248–255, 2009. 1
- [16] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, et al. An image is worth 16x16 words: Transformers for image recognition at scale. In *ICLR*, 2020. 2
- [17] Luciano Floridi and Massimo Chiriatti. Gpt-3: Its nature, scope, limits, and consequences. *Minds and Machines*, 30(4):681–694, 2020. 1
- [18] Jonathan Frankle, Gintare Karolina Dziugaite, Daniel Roy, and Michael Carbin. Linear mode connectivity and the lottery ticket hypothesis. In *ICML*, pages 3259–3269, 2020. 2
- [19] João Gama, Indrè Žliobaitė, Albert Bifet, Mykola Pechenizkiy, and Abdelhamid Bouchachia. A survey on concept drift adaptation. *ACM computing surveys*, 46(4):1–37, 2014. 1
- [20] Xu Han, Zhengyan Zhang, Ning Ding, Yuxian Gu, Xiao Liu, Yuqi Huo, Jiezhong Qiu, Yuan Yao, Ao Zhang, Liang Zhang, et al. Pre-trained models: Past, present and future. *AI Open*, 2:225–250, 2021. 1
- [21] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *CVPR*, pages 770–778, 2016. 1
- [22] Geoffrey Hinton, Oriol Vinyals, and Jeff Dean. Distilling the knowledge in a neural network. *arXiv preprint arXiv:1503.02531*, 2015. 2
- [23] Saihui Hou, Xinyu Pan, Chen Change Loy, Zilei Wang, and Dahua Lin. Learning a unified classifier incrementally via rebalancing. In *CVPR*, pages 831–839, 2019. 2
- [24] Yusong Hu, De Cheng, Dingwen Zhang, Nannan Wang, Tongliang Liu, and Xinbo Gao. Task-aware orthogonal sparse network for exploring shared knowledge in continual learning. In *ICML*, pages 19153–19164, 2024. 2
- [25] Gabriel Ilharco, Marco Tulio Ribeiro, Mitchell Wortsman, Suchin Gururangan, Ludwig Schmidt, Hannaneh Hajishirzi, and Ali Farhadi. Editing models with task arithmetic. In *ICLR*, 2023. 2, 3
- [26] David Isele and Akansel Cosgun. Selective experience replay for lifelong learning. In *AAAI*, pages 3302–3309, 2018. 2
- [27] Menglin Jia, Luming Tang, Bor-Chun Chen, Claire Cardie, Serge J. Belongie, Bharath Hariharan, and Ser-Nam Lim. Visual prompt tuning. In *ECCV*, pages 709–727. Springer, 2022. 1, 2, 3
- [28] 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. *PNAS*, 114(13):3521–3526, 2017. 2
- [29] Chuqiao Li, Zhiwu Huang, Danda Pani Paudel, Yabin Wang, Mohamad Shahbazi, Xiaopeng Hong, and Luc Van Gool. A continual deepfake detection benchmark: Dataset, methods, and essentials. In *WACV*, pages 1339–1349, 2023. 6
- [30] Li Li, Jiawei Peng, Huiyi Chen, Chongyang Gao, and Xu Yang. How to configure good in-context sequence for visual question answering. In *CVPR*, pages 26710–26720, 2024. 1
- [31] Zhizhong Li and Derek Hoiem. Learning without forgetting. *TPAMI*, 40(12):2935–2947, 2017. 2
- [32] Dongze Lian, Zhou Daquan, Jiashi Feng, and Xinchao Wang. Scaling & shifting your features: A new baseline for efficient model tuning. In *NeurIPS*, pages 109–123, 2022. 1

- [33] Yaoyao Liu, Yuting Su, An-An Liu, Bernt Schiele, and Qianru Sun. Mnemonics training: Multi-class incremental learning without forgetting. In *CVPR*, pages 12245–12254, 2020. 2
- [34] Yaoyao Liu, Yingying Li, Bernt Schiele, and Qianru Sun. Wakening past concepts without past data: Class-incremental learning from online placebos. In *WACV*, pages 2226–2235, 2024. 2
- [35] Vincenzo Lomonaco and Davide Maltoni. Core50: a new dataset and benchmark for continuous object recognition. In *Conference on robot learning*, pages 17–26. PMLR, 2017. 6
- [36] 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 *NeurIPS*, 2024. 2
- [37] Marc Masana, Xialei Liu, Bartłomiej Twardowski, Mikel Menta, Andrew D Bagdanov, and Joost van de Weijer. Class-incremental learning: Survey and performance evaluation on image classification. *TPAMI*, 45(05):5513–5533, 2023. 2
- [38] Michael S Matena and Colin A Raffel. Merging models with fisher-weighted averaging. In *NeurIPS*, pages 17703–17716, 2021. 2
- [39] Mark D McDonnell, Dong Gong, Amin Parveneh, Ehsan Abbasnejad, and Anton van den Hengel. Ranpac: Random projections and pre-trained models for continual learning. In *NeurIPS*, 2023. 2, 6
- [40] Martial Mermillod, Aurélia Bugaïska, and Patrick Bonin. The stability-plasticity dilemma: Investigating the continuum from catastrophic forgetting to age-limited learning effects. *Frontiers in psychology*, 4:504, 2013. 1
- [41] Jingyi Ning, Lei Xie, Chuyu Wang, Yanling Bu, Fengyuan Xu, Da-Wei Zhou, Sanglu Lu, and Baoliu Ye. Rf-badge: Vital sign-based authentication via rfid tag array on badges. *IEEE Transactions on Mobile Computing*, 22(02):1170–1184, 2023. 2
- [42] Jingyi Ning, Lei Xie, Yi Li, Yingying Chen, Yanling Bu, Chuyu Wang, Sanglu Lu, and Baoliu Ye. Moirétracker: Continuous camera-to-screen 6-dof pose tracking based on moiré pattern. *IEEE Journal on Selected Areas in Communications*, 42(10):2642–2658, 2024. 1
- [43] Jingyi Ning, Lei Xie, Zhihao Yan, Yanling Bu, and Jun Luo. Moirévision: A generalized moiré-based mechanism for 6-dof motion sensing. In *MobiCom*, pages 467–481, 2024. 1
- [44] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, et al. Pytorch: An imperative style, high-performance deep learning library. In *NeurIPS*, pages 8026–8037, 2019. 6
- [45] Xingchao Peng, Qinxun Bai, Xide Xia, Zijun Huang, Kate Saenko, and Bo Wang. Moment matching for multi-source domain adaptation. In *ICCV*, pages 1406–1415, 2019. 6
- [46] Gabriel Peyré, Marco Cuturi, et al. Computational optimal transport: With applications to data science. *Foundations and Trends® in Machine Learning*, 11(5-6):355–607, 2019. 5
- [47] Jean Piaget. *The construction of reality in the child*. Routledge, 2013. 3
- [48] Zhi-Hong Qi, Da-Wei Zhou, Yiran Yao, Han-Jia Ye, and De-Chuan Zhan. Adaptive adapter routing for long-tailed class-incremental learning. *Machine Learning*, 114(3):1–20, 2025. 2
- [49] Alexandre Ramé, Kartik Ahuja, Jianyu Zhang, Matthieu Cord, Léon Bottou, and David Lopez-Paz. Model ratatouille: Recycling diverse models for out-of-distribution generalization. In *ICML*, pages 28656–28679, 2023. 3
- [50] Roger Ratcliff. Connectionist models of recognition memory: constraints imposed by learning and forgetting functions. *Psychological review*, 97(2):285, 1990. 2, 6
- [51] Sylvestre-Alvise Rebuffi, Alexander Kolesnikov, Georg Sperl, and Christoph H Lampert. icarl: Incremental classifier and representation learning. In *CVPR*, pages 2001–2010, 2017. 2, 6
- [52] David Rolnick, Arun Ahuja, Jonathan Schwarz, Timothy P Lillicrap, and Greg Wayne. Experience replay for continual learning. In *NeurIPS*, pages 350–360, 2019. 2
- [53] Olga Russakovsky, Jia Deng, Hao Su, Jonathan Krause, Sanjeev Satheesh, Sean Ma, Zhiheng Huang, Andrej Karpathy, Aditya Khosla, Michael Bernstein, et al. Imagenet large scale visual recognition challenge. *IJCV*, 115(3):211–252, 2015. 6
- [54] Andrei A Rusu, Neil C Rabinowitz, Guillaume Desjardins, Hubert Soyer, James Kirkpatrick, Koray Kavukcuoglu, Razvan Pascanu, and Raia Hadsell. Progressive neural networks. *arXiv preprint arXiv:1606.04671*, 2016. 2
- [55] Haizhou Shi and Hao Wang. A unified approach to domain incremental learning with memory: Theory and algorithm. In *NeurIPS*, 2023. 2
- [56] Richard Sinkhorn and Paul Knopp. Concerning nonnegative matrices and doubly stochastic matrices. *Pacific Journal of Mathematics*, 21(2):343–348, 1967. 5
- [57] James Seale Smith, Leonid Karlinsky, Vyshnavi Gutta, Paola Cascante-Bonilla, Donghyun Kim, Assaf Arbelle, Rameswar Panda, Rogerio Feris, and Zsolt Kira. Coda-prompt: Continual decomposed attention-based prompting for rehearsal-free continual learning. In *CVPR*, pages 11909–11919, 2023. 2, 3, 6
- [58] Hai-Long Sun, Da-Wei Zhou, Han-Jia Ye, and De-Chuan Zhan. Pilot: A pre-trained model-based continual learning toolbox. *arXiv preprint arXiv:2309.07117*, 2023. 6
- [59] Hai-Long Sun, Da-Wei Zhou, Hanbin Zhao, Le Gan, De-Chuan Zhan, and Han-Jia Ye. Mos: Model surgery for pre-trained model-based class-incremental learning. In *AAAI*, 2025. 2
- [60] Gido M van de Ven, Tinne Tuytelaars, and Andreas S Tolias. Three types of incremental learning. *Nature Machine Intelligence*, pages 1–13, 2022. 1, 2
- [61] Laurens Van der Maaten and Geoffrey Hinton. Visualizing data using t-sne. *JMLR*, 9(11), 2008. 8
- [62] Hemanth Venkateswara, Jose Eusebio, Shayok Chakraborty, and Sethuraman Panchanathan. Deep hashing network for unsupervised domain adaptation. In *CVPR*, pages 5018–5027, 2017. 6
- [63] Cédric Villani. *Optimal transport: old and new*. Springer Science & Business Media, 2008. 5

- [64] Fu-Yun Wang, Da-Wei Zhou, Han-Jia Ye, and De-Chuan Zhan. Foster: Feature boosting and compression for class-incremental learning. In *ECCV*, pages 398–414, 2022. 2
- [65] Fu-Yun Wang, Da-Wei Zhou, Liu Liu, Han-Jia Ye, Yatao Bian, De-Chuan Zhan, and Peilin Zhao. Beef: Bi-compatible class-incremental learning via energy-based expansion and fusion. In *ICLR*, 2023. 2
- [66] 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. 2
- [67] Yabin Wang, Zhiwu Huang, and Xiaopeng Hong. S-prompts learning with pre-trained transformers: An occam’s razor for domain incremental learning. *NeurIPS*, 35:5682–5695, 2022. 1, 2, 3, 6, 7
- [68] 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 *ECCV*, pages 631–648, 2022. 2, 6
- [69] 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 *CVPR*, pages 139–149, 2022. 1, 2, 3, 6, 7
- [70] Max Welling. Herding dynamical weights to learn. In *ICML*, pages 1121–1128, 2009. 7
- [71] Mitchell Wortsman, Gabriel Ilharco, Samir Ya Gadre, Rebecca Roelofs, Raphael Gontijo-Lopes, Ari S Morcos, Hongseok Namkoong, Ali Farhadi, Yair Carmon, Simon Kornblith, et al. Model soups: averaging weights of multiple fine-tuned models improves accuracy without increasing inference time. In *ICML*, pages 23965–23998, 2022. 2
- [72] Mitchell Wortsman, Gabriel Ilharco, Jong Wook Kim, Mike Li, Simon Kornblith, Rebecca Roelofs, Raphael Gontijo Lopes, Hannaneh Hajishirzi, Ali Farhadi, Hongseok Namkoong, et al. Robust fine-tuning of zero-shot models. In *CVPR*, pages 7959–7971, 2022. 2
- [73] Yue Wu, Yinpeng Chen, Lijuan Wang, Yuancheng Ye, Zicheng Liu, Yandong Guo, and Yun Fu. Large scale incremental learning. In *CVPR*, pages 374–382, 2019. 1, 2
- [74] Shipeng Yan, Jiangwei Xie, and Xuming He. Der: Dynamically expandable representation for class incremental learning. In *CVPR*, pages 3014–3023, 2021. 2
- [75] Longrong Yang, Hanbin Zhao, Yunlong Yu, Xiaodong Zeng, and Xi Li. Rcs-prompt: Learning prompt to rearrange class space for prompt-based continual learning. In *ECCV*, pages 1–20, 2024. 2
- [76] Han-Jia Ye, Da-Wei Zhou, Lanqing Hong, Zhenguo Li, Xiu-Shen Wei, and De-Chuan Zhan. Contextualizing meta-learning via learning to decompose. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 46(1):117–133, 2023. 1
- [77] Lu Yu, Bartłomiej Twardowski, Xialei Liu, Luis Herranz, Kai Wang, Yongmei Cheng, Shangling Jui, and Joost van de Weijer. Semantic drift compensation for class-incremental learning. In *CVPR*, pages 6982–6991, 2020. 1
- [78] Friedemann Zenke, Ben Poole, and Surya Ganguli. Continual learning through synaptic intelligence. In *ICML*, pages 3987–3995, 2017. 2
- [79] Jianyu Zhang and Léon Bottou. Learning useful representations for shifting tasks and distributions. In *ICML*, pages 40830–40850, 2023. 3
- [80] Bowen Zhao, Xi Xiao, Guojun Gan, Bin Zhang, and Shu-Tao Xia. Maintaining discrimination and fairness in class incremental learning. In *CVPR*, pages 13208–13217, 2020. 1, 2
- [81] Wenyi Zhao, Rama Chellappa, P Jonathon Phillips, and Azriel Rosenfeld. Face recognition: A literature survey. *ACM computing surveys*, 35(4):399–458, 2003. 1
- [82] Bowen Zheng, Da-Wei Zhou, Han-Jia Ye, and De-Chuan Zhan. Multi-layer rehearsal feature augmentation for class-incremental learning. In *ICML*, pages 61649–61663, 2024. 2
- [83] Bowen Zheng, Da-Wei Zhou, Han-Jia Ye, and De-Chuan Zhan. Task-agnostic guided feature expansion for class-incremental learning. In *CVPR*, 2025. 2
- [84] Da-Wei Zhou, Qi-Wei Wang, Han-Jia Ye, and De-Chuan Zhan. A model or 603 exemplars: Towards memory-efficient class-incremental learning. In *ICLR*, 2023. 2, 6
- [85] Da-Wei Zhou, Hai-Long Sun, Han-Jia Ye, and De-Chuan Zhan. Expandable subspace ensemble for pre-trained model-based class-incremental learning. In *CVPR*, 2024. 6
- [86] Da-Wei Zhou, Qi-Wei Wang, Zhi-Hong Qi, Han-Jia Ye, De-Chuan Zhan, and Ziwei Liu. Class-incremental learning: A survey. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 46(12):9851–9873, 2024. 2
- [87] Da-Wei Zhou, Zi-Wen Cai, Han-Jia Ye, De-Chuan Zhan, and Ziwei Liu. Revisiting class-incremental learning with pre-trained models: Generalizability and adaptivity are all you need. *International Journal of Computer Vision*, 133:1012–1032, 2025. 2, 6
- [88] Fei Zhu, Xu-Yao Zhang, Chuang Wang, Fei Yin, and Cheng-Lin Liu. Prototype augmentation and self-supervision for incremental learning. In *CVPR*, pages 5871–5880, 2021. 2