

Towards Adapting Vision-Language Models for Semi-Supervised Domain Generalization

Xilin He*, Xiaole Xian[†], Zhihua Liu[‡], Weicheng Xie[†], Xiangyu Yue[§], Muhammad Haris Khan*

*Mohamed bin Zayed University of Artificial Intelligence (MBZUAI)

[†]Shenzhen University

[‡]University of Edinburgh

[§]The Chinese University of Hong Kong

Corresponding author: muhammad.haris@mbzuai.ac.ae

Abstract—Semi-supervised Domain Generalization (SSDG) offers a cost-effective solution for generalizing models to unseen domains with limited labels. While existing SSDG methods, mainly built upon small-scale backbones, struggle to match fully supervised DG performance, large-scale vision-language models like CLIP have shown remarkable generalization through downstream fine-tuning. However, adapting these models to SSDG remains underexplored. In this paper, we identify a critical issue: existing popular fine-tuning methods suffer from under-utilizing unlabeled data in the semi-supervised learning frameworks, thereby overfitting the limited labeled data, leading to training collapse and generalization ability degradation. To address these challenges, we propose two novel components: (1) the De-False-Correlation Adapter (DFC-Adapter), which reduces false correlations to refine visual features and (2) Learnable Multi-granularity Text-guided Embedding Augmentation (LMTEA), which synthesizes semantic-aligned but domain-perturbed augmented visual embedding for consistency regularization through multi-granularity text guidance and learnable style encoding. Moreover, we establish the first-ever benchmark for CLIP fine-tuning methods in SSDG, conducting extensive experiments across six DG datasets and two ImageNet variants. Our results demonstrate that our method significantly outperforms existing CLIP fine-tuning approaches and achieves performance comparable to even fully supervised DG methods in some cases.

I. INTRODUCTION

Semi-supervised domain generalization (SSDG) aims to learn models that generalize to unseen domains by leveraging both limited labeled and abundant unlabeled data from multiple source domains. As the existing methods mainly build upon FixMatch [1] and its variants [2], [3], [4], [5], they focus on three major technical directions of improving pseudo-labeling (PL) accuracy [2], [6], [7], [8], modeling domain-level information [9], [10] and data-level consistency regularization [4]. However, we notice the existence of three fundamental limitations: (1) Architectural constraints. Current approaches rely on small-scale backbones (e.g., ResNet-18 [11]), which lack the scalability and generalization capacity of modern vision-language models (VLMs). (2) Pseudo-label reliability. Low PL accuracy injects noise into training, propagating errors and degrading generalization, which is a critical weakness given the limited labeled data in SSDG. (3) Augmentation poverty. Existing works [12], [9], [4], [5] mainly use simple predefined image-level augmentations, such as image rotation and flipping, hindering robustness to diverse

domain shifts. Collectively, these limitations constrain further improvement of SSDG, precluding competitive performance against fully supervised DG methods.

Recent advances in adapting VLMs like CLIP [13] through downstream fine-tuning [14], [15], [16] offer new opportunities. Popular fine-tuning methods, such as low-rank adaptation (LoRA) [14] and prompt tuning techniques [15], [16], offer promising pathways for customizing foundational models while preserving transferable knowledge. However, their application in SSDG scenarios remains underexplored. A critical limitation emerges here: existing downstream fine-tuning methods severely underutilize unlabeled data during SSDG training. As illustrated in Figure. 1, pseudo-labeling confidence from CLIP’s image-text pairing often falls below the hand-crafted confidence threshold, rendering large portions of unlabeled data inactive in gradient updates under the semi-supervised learning framework [1]. This exacerbates the confirmation bias [17] inherent in semi-supervised learning, leading to overfitting on the limited labeled data and degrading generalization performance. Despite LoRA achieving high PL accuracy in the initial training period, only a small portion of unlabeled data contributes to gradient updates. Meanwhile, directly applying linear probing achieves marginal performance, but still degrades the pseudo-labeling accuracy due to potential overfitting. Another straightforward alternative, combining fine-tuning methods with a linear classifier, leads to training collapses similar to those of the vanilla fine-tuning. This suggests that preventing learnable modules from underutilizing the unlabeled data and overfitting to the limited labeled data is crucial for adapting VLMs to SSDG.

To address the challenges above and bridge the gap in VLMs research in SSDG, this paper proposes a De-False-Correlation Adapter (DFC-Adapter) and Learnable Multi-granularity Text-guided Embedding Augmentation (LMTEA) to prevent the confirmation bias from the perspective of architecture design and data augmentation. More specifically, DFC-Adapter learns generalizable knowledge to refine visual features in both spatial and semantic space by decreasing false correlations from both domain-specific biases and pre-trained knowledge. Meanwhile, LMTEA achieves a richer augmentation space for consistency learning with learnable style encoding and text-guided embedding augmentation from both object attribute

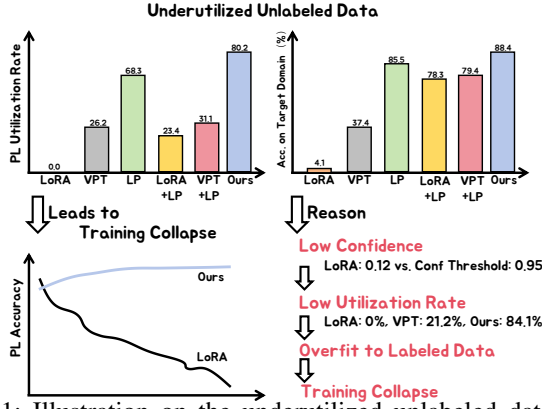


Fig. 1: Illustration on the underutilized unlabeled data issue of existing fine-tuning methods (LoRA [14] and VPT [15]) under SSDG setting. LP is linear probing.

level and global style level. We evaluate our method on standard DG benchmark datasets and ImageNet variants, establishing the first benchmark for VLM fine-tuning methods in SSDG. Experimental results demonstrate that our method significantly outperforms existing methods and, in some cases, achieves performances comparable to fully supervised DG baselines. Our main contributions can be summarized as:

- We identify the low unlabeled data utilization in prior VLM fine-tuning methods under the SSDG scenarios, leading to training collapse and degradation of generalization performances.
- From the perspective of architecture design and data augmentation, we propose DFC-Adapter and LMTEA to adapt pre-trained VLMs to SSDG by preventing overfitting to the limited label data and training collapses caused by confirmation bias.
- Extensive experimental results across 8 datasets, where we achieve comparative performances with the fully supervised methods, demonstrate the superiority of the proposed method.

II. RELATED WORK

Domain Generalization: Domain generalization (DG) intends to learn a model from (multiple) source domains with transferable knowledge that can generalize to previously unseen target domains. Existing methods could be substantially categorized into data augmentation [18], [19], [20], [21], [22], domain alignment [23], [24], meta-learning [25], [26] and optimization methods [27], [28], [29]. Despite gaining performance improvements, the vast majority of the DG research is ill-equipped to process unlabeled data, largely stemming from their foundational assumption of a fully supervised learning context. However, in real-world scenarios, it could be infeasible to curate a fully-labeled dataset with various domains for training, thereby limiting the applications of these fully supervised DG methods.

Semi-supervised Domain Generalization: Semi-supervised domain generalization (SSDG) has emerged as a promising avenue to address domain shift with limited labeled data. Existing SSDG methods [1], [4], [7], [9], [12] are mainly

developed from semi-supervised learning methods, such as FixMatch [1], demonstrating the crucial impact of pseudo-labeling (PL) and consistency regularization on SSDG performances. A handful of research [12], [9], [7], [10] has been dedicated to improving PL accuracy during training. [10] approach domain-aware PL with a dual-classifier structure, while [9] directly modulates the weight matrix of the classifier head with domain-level information. Meanwhile, on the aspect of consistency regularization, StyleMatch [4] introduces stochastic modeling on the classifier head and MultiMatch [5] formulates the multiple source domains into multiple local tasks and a global task for domain alignment. Despite demonstrating notable performance improvements, these methods significantly trail fully supervised DG performances. To the best of our knowledge, the vast majority of SSDG research exclusively applies small-scale convolutional networks as their backbones, limiting their further development. Meanwhile, for consistency learning, these methods depend on predefined image-level augmentation to generate augmented views in semi-supervised learning, which overlooks the rich semantic augmentation space.

VLM Generalization: Foundation vision-language models like CLIP [30], pre-trained on web-scale data, achieve strong out-of-distribution generation [31]. However, direct fine-tuning with task-specific data often harms robustness to distribution shifts [32]. To address this, parameter-efficient strategies, such as LoRA [14], CoOP [16], VPT [15], and MaPLE [33] adapt only a few parameters while preserving generalizable knowledge. Beyond these, advanced methods further enhance OOD generalization via style-aware prompting [34], disentangled representations [35], or knowledge distillation [36]. Despite these efforts, most fine-tuning methods assume fully supervised settings, leaving their behavior in semi-supervised domain generalization underexplored. In parallel, feature augmentation approaches [37], [38] enrich the training distribution by synthesizing semantically aligned but domain-perturbed features, enabling consistency regularization across domains. However, building upon the modality gap assumption [39], they assume predefined text based on the domain name would be a perfect match of the visual features, which is not the case in real-world scenarios, thereby leading to potential semantically misaligned augmented features.

III. METHOD

A. Preliminaries

Problem Settings. We denote each domain d by $d = \{(x_i^d, y_i^d)\}_{i=1}^n$, where x_i^d , y_i^d and n is an input image, the corresponding ground-truth label and the total number of images in domain d , respectively. In the scenario of SSDG, there are only limited labeled samples, each source domain has a labeled part $d^L = \{(x_i^d, y_i^d)\}$ and an unlabeled part $d^U = \{(u_i^d)\}$, where the number of samples in the unlabeled part is significantly higher than that in the labeled part. For the vision-language model CLIP [30], we denote its image encoder and text encoder as E_v and E_t , respectively.

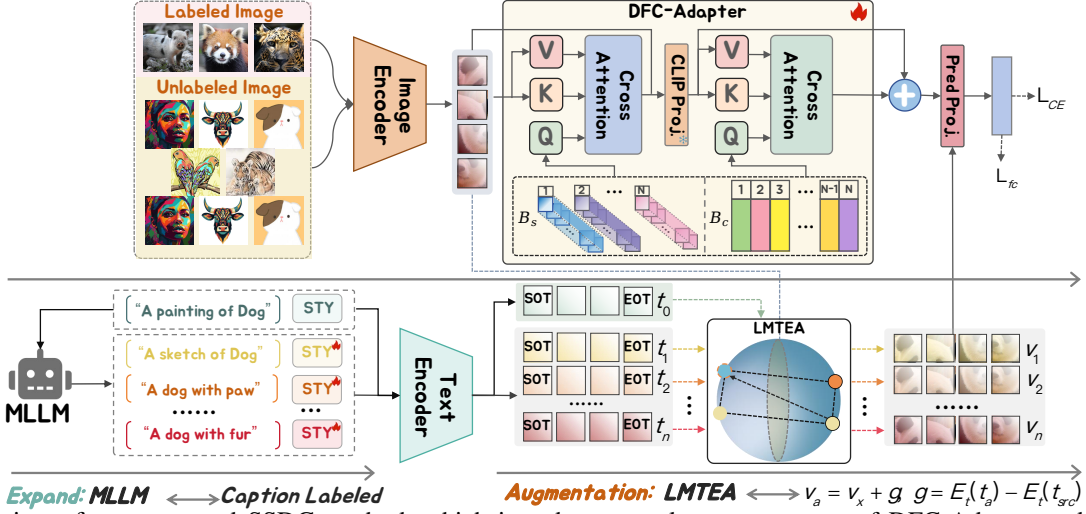


Fig. 2: Overview of our proposed SSDG method, which introduces two key components of DFC-Adapter and LMTEA. B_s and B_c denote the spatial refinement bank and vision-language alignment correlation bank, respectively.

SSDG Pipeline. We adopt FixMatch [1] as our baseline due to its empirical effectiveness in SSDG and conceptual simplicity. It integrates two SSL mechanisms. Pseudo-labeling (PL) assigns labels to unlabeled data when their maximum class probability exceeds a confidence threshold. Meanwhile, consistency regularization enforces prediction invariance across augmented views. FixMatch processes each sample in a mini-batch with both weak and strong augmentations. The total loss consists of: 1) supervised loss $\mathcal{L}_s$ applied on the weakly augmented version of the labeled data. And 2) $\mathcal{L}_u$ aligns predictions on strongly augmented unlabeled data with their pseudo-labels (generated from the weakly augmented versions). Despite its success in SSDG, FixMatch [1] faces significant challenges when adapted to downstream tasks with VLMs. As illustrated in Figure. 1, prevalent fine-tuning approaches fail to effectively utilize unlabeled data within FixMatch’s pseudo-labeling framework, resulting in both training collapse and generalization degradation.

B. Method Overview

As illustrated in Figure. 2, our framework comprises two key components: De-False-Correlation Adapter (DFC-Adapter) and Learnable Multi-granularity Text-guided Embedding Augmentation (LMTEA). Given visual features extracted by the image encoder E_v , the DFC-Adapter applies learnable knowledge banks as corrections to mitigate both domain-specific information and spurious correlations. During training, LMTEA synthesizes augmentation embeddings in the vision-language embedding space by incorporating learnable style encoding and external knowledge from large-language models (LLMs) to enhance consistency regularization.

C. De-False-Correlation Adapter Learning

CLIP’s pretrained visual encoder E_v often encodes spurious correlations between specific components, stemming from the source domains or pretrained knowledge, and semantic content [40], leading to generalization degradation. To mitigate

this while preserving cross-modal alignment and prevent overfitting, we propose the De-False-Correlation Adapter (DFC-Adapter) P , which dynamically suppresses biased feature activations through two sets of learnable knowledge bank: (1) a spatial refinement bank $B_s = \{b_{s,k}\}_{k=1}^K \in \mathbb{R}^{K \times d_s}$ that refines visual features over spatial tokens to enhance discriminative local patterns and (2) a semantic correlation alignment bank $B_c = \{b_{c,k}\}_{k=1}^K \in \mathbb{R}^{K \times d_c}$. K , d_s and d_c denote the number of learnable features in each bank and their corresponding feature dimensions, respectively. For a batch of input visual features $z = E_v(x) \in \mathbb{R}^{B \times L \times d_s}$, we first compute cross-attention between B_s and z to adaptively adjust spatial feature activations:

$$\begin{cases} P(z) = W_P(CA(z, b_s)), & CA(z, b_s) = \text{softmax}(\frac{Q_s K_s^T}{\sqrt{d}}) V_s \\ Q_s = W_{Q,s} b_s, & K_s = W_{K,s} z, & V_s = W_{V,s} z, \end{cases} \quad (1)$$

where $W_{Q,s}$, $W_{K,s}$ and $W_{V,s}$ are the corresponding attention mapping matrix. W_P denotes the mapping matrix for dimension alignment in the DFC-Adapter P . Then, we combine the output with the original visual feature in the vision-language alignment space:

$$z_s = \text{CLS}(z) + P(z), \quad (2)$$

where $\text{CLS}(\cdot)$ denotes projecting the class token to vision-language space. z_s denotes the semantic visual feature in the vision-language embedding space. Subsequently, we incorporate the semantic correlation bank B_c with the projected feature z_s via cross-attention mechanism similar with that in Eq. 1 as:

$$z_f = z_f + CA(z_f, b_c), \quad (3)$$

where z_f denotes the final feature for linear class prediction. **De-False-Correlation Learning.** To mitigate domain-specific and pretrained knowledge biases while preserving semantic fidelity and cross-modality alignment, we propose a dual-level regularization mechanism. From the domain level, to decrease domain-specific bias information in the feature, the output of the DFC-Adapter should be varied across inputs from different domains and should be similar when inputs from the same

source domain are given. Therefore, we formulate this as a contrastive domain alignment loss:

$$\mathcal{L}_{da} = -\frac{1}{n} \sum_{i=1}^n \log \frac{\sum_{i \neq j, d(z_i)=d(z_j), j=1, \dots, n} e^{-D(P(z_i), P(z_j))}}{\sum_{i \neq k, k=1, \dots, n} e^{-D(P(z_i), P(z_k))}}, \quad (4)$$

where n , $d(\cdot)$ and $D(\cdot)$ denote the batch size, obtaining the domain index from the source domains and L2 distance, respectively.

From the object-centric level, to mitigate spurious false correlations from the visual encoder of CLIP (deeming ‘chair’ and ‘desk’ identical), we propose to leverage the external knowledge from existing large-language models (LLMs). We prompt LLMs with a template to list commonly objects co-occurring with but essentially irrelevant to each class c_{src} . For example, $c_{fc, chair} = \{\text{desk, office, ...}\}$. Based on these counterfactual objects, we pull the visual feature of c_{src} away from those text embeddings of c_{fc} , while maintaining image-text alignment with the original class text embedding as:

$$\mathcal{L}_{fc} = -\frac{1}{n} \sum_{i=1}^n \left(\sum_{j=1}^{|c_{fc, y_i}|} z_{f, i, j} E_t(c_{fc, y_j}) - \text{ITP}(z_{f, i, j}, E_t(c_{src})) \right), \quad (5)$$

where c_{fc, y_i} denote the false-correlation text list corresponding to class of y_i . ITP denotes calculating the cross-entropy loss on the image-text pairing prediction with the ground-truth label. The loss calculation object under the semi-supervised learning framework will be illustrated in detail in the latter section.

D. Learnable Multi-granularity Text-guided Embedding Augmentation

To further enhance consistent regularization, we propose Learnable Multi-granularity Text-guided Embedding Augmentation (LMETA), which synthesizes semantic-aligned and domain-perturbed visual embeddings as augmentation samples during training. Unlike existing methods of VLM feature augmentation that rely on pre-defined domain names, LMETA introduces learnable style encoding and multi-granularity text-guided augmentation, effectively addressing semantic misalignment and enhancing the diversity of augmented features. **Analysis on TEAM [38].** TEAM exploits the modality gap phenomenon [39], [41], which is the constant vector orthogonal to the span of image and text embeddings, to translate original visual features into novel domains. As shown in Figure. 3(a), the TEAM directly assumes that the text t_{src} in the format of “a {source domain name} of the {class name}” would be the one match for the images from the source domain, and synthesize visual features from novel domains with another text t_a “a {augmented domain name} of the {class name}” by:

$$v_a = v_x + g, \quad g = E_t(t_a) - E_t(t_{src}), \quad v_x = E_v(x) \quad (6)$$

However, the effectiveness of TEAM is based on the assumption that the predefined texts would be a well-matched pair for the image features. It fails when the domain name fails to represent the domain-specific information (domain name is set as the number of camera traps in TerraIncognita [42]), leading to potential semantic misaligned augmented embedding and hindering consistency learning.

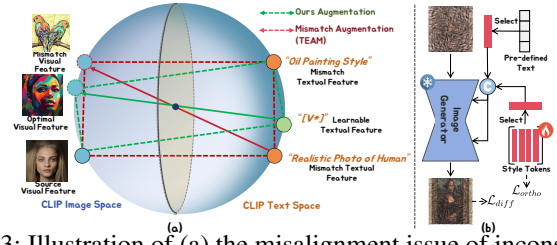


Fig. 3: Illustration of (a) the misalignment issue of inconsistent modality gap from TEAM [38], and (b) the training process of the learnable style encoding in LMTEA.

Learnable Style Encoding. To overcome this limitation, we propose to use the learnable style embedding to replace source domain name embedding. Inspired by previous work on image editing with diffusion models [43], [44], where a learnable component is integrated with text embedding to generate novel concepts, we incorporate a set of learnable style components s . As shown in Figure. 3(b), given an image x from the d -th domain, we select the corresponding style components to combine with the pre-defined text for noise prediction with the diffusion models. We optimize the learnable style component with the diffusion loss [45] to learn delicate domain-specific style components. To further improve the learning of domain-specific information, we apply an orthogonality loss to reduce the correlation between each pair of style components:

$$\mathcal{L}_{ortho} = \|ss^T - \text{diag}(ss^T)\|_2, \quad (7)$$

where diag means keeping only the diagonal entries. We combine the above two losses to learn the domain-specific style component before launching the fine-tuning of CLIP. With the learned style encoding, we combine them with the pre-defined text embedding to generate augmentation embeddings:

$$x_{aug} = v_x + \hat{g}, \quad \hat{g} = E_t(t_a) - \text{concat}(E_t(t_{src}) + s_d), \quad (8)$$

where concat means concatenate the text embeddings and d denotes the domain index for selecting the corresponding style component.

Multi-granularity Text-guided Augmentation. To facilitate generalized learning at the object level, we propose a text-guided multi-granularity augmentation strategy based on Eq. 8, including both style-level and attribute-level augmentation. For style-level shifting, we prompt an LLM to generate a set of style words t_{sty} and prepare the augmented text in the format of “a {style word} of {class name}”. Similarly to style-level shifting, for attribute-level augmentation, we prompt the LLM to generate the attribute sets for each class and craft the augmented text as “a {class name} with {attribute}”.

E. Training Loss

Before fine-tuning the CLIP, we first train the domain-specific style encoding with a diffusion-based image generation model, as illustrated in the LMTEA section. Subsequently, we proceed to fine-tune the CLIP model with the DFC-Adapter. In the semi-supervised framework with both labeled and unlabeled data, the overall training loss is divided

into supervised and unsupervised branches, following the FixMatch [1].

For labeled data, we compute cross-entropy losses for both original and augmented embeddings crafted by the proposed multi-granularity text-guided augmentation. And Jensen-Shannon divergence is applied to further enhance those embeddings consistency. Meanwhile, the domain-aware loss in Eq. 4, the false-correlation regularization in Eq. 5 are applied to the supervised data. Thereby, the supervised branch of the training loss is formulated as:

$$\mathcal{L}_{sup} = \mathcal{L}_{CE}(\hat{y}, y) + \mathcal{L}_{CE}(\hat{y}_{aug}, y) + \lambda_{JS} \mathcal{L}_{JS}(\hat{y}, \hat{y}_{aug}) + \lambda_{auxsup}(\mathcal{L}_{fc} + \mathcal{L}_{da}), \quad (9)$$

where $\mathcal{L}_{CE}$, $\hat{y}$ and $\hat{y}_{aug}$ denote the cross-entropy loss, prediction on the original visual embedding and the augmented embedding of the labeled data, respectively.

For the unlabeled data, we first generate pseudo label for the unlabeled images by confidence threshold, obtaining a mask m to select samples with high confidence predictions for loss calculation. A cross-entropy loss is then computed between the predictions of weakly and strongly augmented views, while a Jensen-Shannon divergence is applied between the weakly augmented view and the augmented embedding for the selected samples. Therefore, the unsupervised loss is formulated as:

$$\mathcal{L}_{unsup} = m(\mathcal{L}_{CE}(\hat{y}_w, \hat{y}_s) + \lambda_{JSu} \mathcal{L}_{JSu}(\hat{y}_w, \hat{y}_{aug})), \quad (10)$$

where $\hat{y}_w$ and $\hat{y}_s$ denote the predictions for the weakly and strongly augmented view. $\mathcal{L}_{JSu}$ represents Jensen-Shannon divergence loss calculated on the unsupervised data. m is the mask for high confidence sample selection. Finally, the model is trained with the combination of $\mathcal{L}_{sup}$ and $\mathcal{L}_{unsup}$.

IV. EXPERIMENTS

A. Datasets, Settings and Implementation Details

Datasets. We evaluate on standard DG benchmarks: PACS [46], Office-Home [47], Digits [48], TerraIncognita [42], and VLCS [49], plus the large-scale DG dataset DomainNet [50]. For scalability tests, we use semi-supervised ImageNet [51] for training and evaluate on corrupted variants (ImageNet-A [52] and ImageNet-R [53]).

Settings. Two established SSDG settings [12], [9], [4], [10] are evaluated. (1) 5 or 10 labeled samples per class per source domain [12], [9]. (2) One source domain is fully labeled, while the others are unlabeled [4], [10]. The above two settings are referred to as the 1st and 2nd settings in the following text, respectively. For ImageNet-scale experiments, we fine-tune pretrained models with 1% labeled data. More details can be found in the supplementary material.

Evaluation Protocol. We follow the leave-one-domain-out protocol for evaluation. Following [12], [9], we report the average performances over 5 independent runs.

Baselines. We evaluate with: (1) standard downstream fine-tuning methods, including LoRA [14], VPT [15], CLIPood [31], CoOp [16], PromptSRC [54]. (2) SOTA SSDG

method UPCSC [55] and DGWM [9]. (3) CLIP fine-tuning methods for supervised DG, such as VL2V-SD [36] and MoA [56]. All methods are integrated with FixMatch [1]. For reference, we also report fully supervised results of CLIP fine-tuning methods.

B. Main Results

Results on standard DG benchmarks. Tables. I and II demonstrate our method’s superior performances across DG benchmarks. We significantly outperform fine-tuning methods (LoRA [14], CoOp [16], VPT [15]), which show substantial degradation when adapted to SSDG settings due to limited unlabeled data utilization. Notably, our approach exceeds robust fine-tuning baseline MoA [56] by 2.28% and 2.38% on TerraIncognita in Setting 1 and rivals fully supervised SOTA methods on PACS/OfficeHome. Further surpassing SSDG SOTA DGWM [9], we achieve 0.87% and 0.85% gains on DomainNet under both settings.

Results on ImageNet variants. Table. IIIa validates the efficiency of the proposed method on corrupted ImageNet variants, where we outperform the previous robust fine-tuning SOTA MoA [56] by 1.44% and 5.65% on ImageNet-A [52] and ImageNet-R [57], respectively. This pronounced improvement under ImageNet-R’s challenging distribution shifts confirms the superiority of LMTEA in simulating visual embeddings from different unseen domains.

C. Ablation Study

Contributions of each component. We evaluate key component contributions through ablation studies in Table. IIIb. As shown, the DFC-Adapter alone improves linear probing performances by 2.54%, demonstrating its effectiveness in refining representations. Conversely, removing our learnable style encoding causes significant degradation, confirming its essential role in generating semantically-aligned augmented embeddings for domain generalization.

Training evolution. We present the evolution of the PL utilization rate and PL accuracy between LoRA [14] and our approach throughout the training process in Figure. 4(a). As LoRA continually fails to adapt more unlabeled data for training, it overfits the limited labeled data, leading to degradation of the PL accuracy. However, after the initial training period, our method uses a growing number of unlabeled data for training, leading to high PL accuracy and effective training.

Hyperparameter Studies. We conduct hyperparameter analysis on the size of the knowledge banks, where K_s and K_c define the size of the spatial refinement bank and semantic alignment bank, as well as the weights of the auxiliary loss weight λ_{auxsup} , on the Office-Home dataset [47]. As shown in Figure. 4(b), when we set a moderate size for the learnable banks, the performances remain stable overall, indicating that our method is not sensitive to the bank size generally. However, with an excessive bank size, the performances undergo significant degradation, as the bank with such a size would not be trained thoroughly and inject potential noise during inference. In terms of the weight for the auxiliary loss, the

Method			5 labels per class						10 labels per class			
	PACS	OfficeHome	Digits	TerraInc.	VLCS	DomainNet	PACS	OfficeHome	Digits	TerraInc.	VLCS	DomainNet
<i>Fully Supervised</i>												
STYLIP	98.05	84.63	81.38	-	86.94	62.02	98.05	84.63	81.38	-	86.94	62.02
CoOp	97.00	81.12	76.41	-	82.98	59.52	97.00	81.12	76.41	-	82.98	59.52
VPT	96.90	83.20	-	46.70	82.00	58.50	96.90	83.20	-	46.70	82.00	58.50
<i>Semi-Supervised</i>												
FixMatch												
+Linear Probe	96.27	83.95	62.14	34.78	72.62	58.43	96.18	84.68	62.31	35.43	72.31	59.67
+LoRA	95.28	78.19	60.37	15.24	75.82	52.32	95.72	78.32	60.47	15.63	76.27	53.51
+CoOP	96.25	81.39	62.47	33.98	74.08	54.62	96.31	81.55	62.62	35.12	74.89	54.98
+VPT	96.19	80.33	61.86	35.64	76.42	49.81	96.23	80.54	62.43	35.56	76.53	50.42
+CLIPood	95.44	76.43	62.41	35.90	75.11	56.42	96.19	76.99	62.95	36.56	75.58	56.67
+PromptSRC	96.38	81.59	63.19	35.43	76.55	-	96.47	82.42	63.50	35.98	76.76	-
+VL2V-SD	95.46	85.47	66.40	38.42	76.74	57.92	95.68	86.12	66.98	38.96	77.83	58.41
+MoA	96.23	87.62	73.42	39.16	77.05	59.62	96.59	88.23	74.14	39.85	77.68	59.78
+DGWM	96.54	85.47	63.80	36.23	73.99	58.72	96.59	85.94	64.04	36.59	74.24	58.98
+UPCSC	96.68	86.32	69.84	37.04	77.49	-	96.78	87.21	73.66	38.48	78.25	-
+Ours	96.74	88.68	76.24	41.44	78.17	59.59	96.83	88.94	75.83	42.23	78.12	59.83

TABLE I: Comparison with fine-tuning methods and SSDG SOTAs on popular DG benchmark datasets under the first setting. The best results under semi-supervised frameworks are in bold.

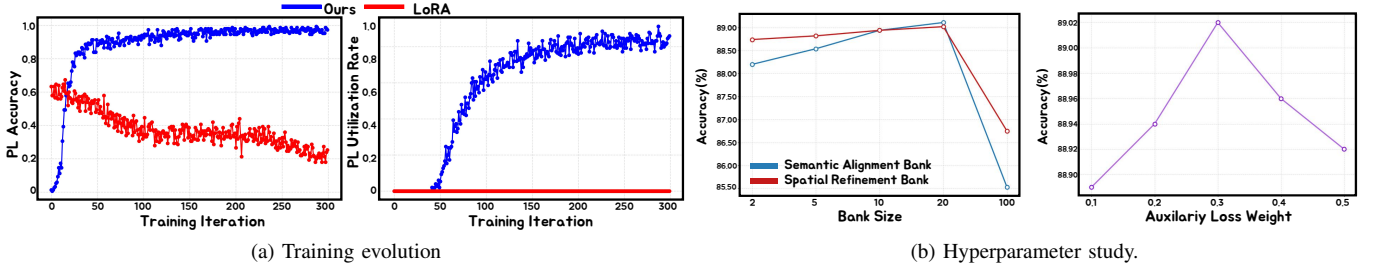


Fig. 4: (a) Evolution of the pseudo-label (PL) accuracy and unlabeled data utilization rate during training on OfficeHome [47]. (b) Hyperparameter study on the learnable banks' sizes in DFC-Adapter and the auxiliary losses weight λ_{auxsup} .

Method (FixMatch)	PACS	OfficeHome	Digits	TerraInc.	VLCS	DomainNet
+Linear Probe	96.41	84.52	57.04	32.69	66.05	49.79
+LoRA	88.59	73.24	52.88	14.03	62.31	48.34
+CoOP	92.46	82.43	56.70	29.35	58.94	49.05
+VPT	89.93	74.90	54.38	12.53	64.83	47.90
+CLIPood	89.25	75.21	55.63	28.04	63.13	51.43
+PromptSRC	94.24	79.42	57.84	31.53	64.08	-
+MoA	96.84	85.19	59.37	34.65	68.93	54.78
+DGWM	96.44	84.98	58.14	35.68	68.51	54.65
+Ours	96.89	86.73	64.89	39.59	72.52	54.88

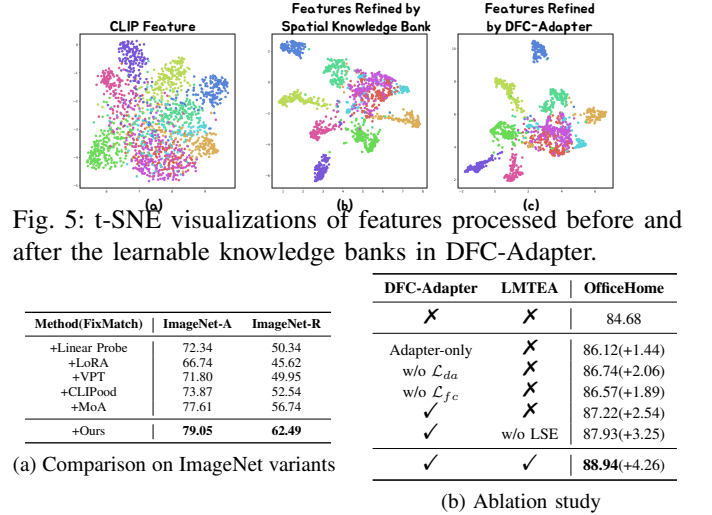
TABLE II: Comparison on popular DG benchmark datasets under the second SSDG setting. The best results are labeled in bold.

changes in the weight merely affect the performance in a small margin.

Effectiveness of feature refinement from the Learnable Knowledge Banks. To validate the effectiveness of the learnable knowledge banks on feature refinement, we provide t-SNE visualizations of the process of the feature changes when passing through the DFC-Adapter. As shown in Figure. 5, after interacting with the two sets of knowledge banks, the features become better clustered and more discriminative than the original CLIP features, demonstrating that the learned knowledge banks can refine the visual features for better performance.

V. CONCLUSION

In this paper, we address the challenge of adapting vision-language models (VLMs) for semi-supervised domain generalization (SSDG). We reveal that existing downstream fine-tuning methods suffer from low utilization rates of unlabeled



(a) Comparison on ImageNet variants

TABLE III: Comparison on ImageNet variants and ablation study.

data, leading to overfitting and degraded generalization. To tackle this, we propose DFC-Adapter and LMTEA, which aim to prevent the confirmation bias of semi-supervised learning from both perspectives of architecture design and data augmentation. Our method significantly outperforms prior methods, achieving comparative results to fully supervised methods on several datasets. This work establishes the first benchmark for VLM SSDG, highlighting the potential of adapting VLMs for robust generalization under limited labeled data.

REFERENCES

- [1] K. Sohn, D. Berthelot, N. Carlini, Z. Zhang, H. Zhang, C. A. Raffel, E. D. Cubuk, A. Kurakin, and C.-L. Li, “Fixmatch: Simplifying semi-supervised learning with consistency and confidence,” *Advances in neural information processing systems*, vol. 33, pp. 596–608, 2020. 1, 2, 3, 5
- [2] B. Zhang, Y. Wang, W. Hou, H. Wu, J. Wang, M. Okumura, and T. Shinozaki, “Flexmatch: Boosting semi-supervised learning with curriculum pseudo labeling,” *Advances in Neural Information Processing Systems*, vol. 34, pp. 18408–18419, 2021. 1
- [3] Y. Wang, H. Chen, Q. Heng, W. Hou, Y. Fan, Z. Wu, J. Wang, M. Savvides, T. Shinozaki, B. Raj *et al.*, “Freematch: Self-adaptive thresholding for semi-supervised learning,” *arXiv preprint arXiv:2205.07246*, 2022. 1
- [4] K. Zhou, C. C. Loy, and Z. Liu, “Semi-supervised domain generalization with stochastic stylematch,” *International Journal of Computer Vision*, vol. 131, no. 9, pp. 2377–2387, 2023. 1, 2, 5
- [5] L. Qi, H. Yang, Y. Shi, and X. Geng, “Multimatch: Multi-task learning for semi-supervised domain generalization,” *ACM Transactions on Multimedia Computing, Communications and Applications*, vol. 20, no. 6, pp. 1–21, 2024. 1, 2
- [6] S. Zoha, J.-G. Lee, and Y.-W. Ko, “Cat: Class aware adaptive thresholding for semi-supervised domain generalization,” *arXiv preprint arXiv:2412.08479*, 2024. 1
- [7] A. Khan, M. A. Shaaban, and M. H. Khan, “Improving pseudo-labelling and enhancing robustness for semi-supervised domain generalization,” *arXiv preprint arXiv:2401.13965*, 2024. 1, 2
- [8] L. Qi, H. Yang, Y. Shi, and X. Geng, “Multimatch: Multi-task learning for semi-supervised domain generalization,” *ACM Transactions on Multimedia Computing, Communications and Applications*, vol. 20, no. 6, pp. 1–21, 2024. 1
- [9] C. J. Galappaththige, Z. Izzo, X. He, H. Zhou, and M. H. Khan, “Domain-guided weight modulation for semi-supervised domain generalization,” *arXiv preprint arXiv:2409.03509*, 2024. 1, 2, 5
- [10] R. Wang, L. Qi, Y. Shi, and Y. Gao, “Better pseudo-label: Joint domain-aware label and dual-classifier for semi-supervised domain generalization,” *Pattern Recognition*, vol. 133, p. 108987, 2023. 1, 2, 5
- [11] K. He, X. Zhang, S. Ren, and J. Sun, “Deep residual learning for image recognition,” in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2016, pp. 770–778. 1
- [12] C. J. Galappaththige, S. Baliah, M. Gunawardhana, and M. H. Khan, “Towards generalizing to unseen domains with few labels,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024, pp. 23 691–23 700. 1, 2, 5
- [13] P. Gao, S. Geng, R. Zhang, T. Ma, R. Fang, Y. Zhang, H. Li, and Y. Qiao, “Clip-adaptor: Better vision-language models with feature adapters,” *International Journal of Computer Vision*, vol. 132, no. 2, pp. 581–595, 2024. 1
- [14] E. J. Hu, Y. Shen, P. Wallis, Z. Allen-Zhu, Y. Li, S. Wang, L. Wang, W. Chen *et al.*, “Lora: Low-rank adaptation of large language models,” *ICLR*, vol. 1, no. 2, p. 3, 2022. 1, 2, 5
- [15] M. Jia, L. Tang, B.-C. Chen, C. Cardie, S. Belongie, B. Hariharan, and S.-N. Lim, “Visual prompt tuning,” in *European conference on computer vision*. Springer, 2022, pp. 709–727. 1, 2, 5
- [16] K. Zhou, J. Yang, C. C. Loy, and Z. Liu, “Conditional prompt learning for vision-language models,” in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2022, pp. 16816–16825. 1, 2, 5
- [17] E. Arazo, D. Ortego, P. Albert, N. E. O’Connor, and K. McGuinness, “Pseudo-labeling and confirmation bias in deep semi-supervised learning,” in *2020 International Joint Conference on Neural Networks (IJCNN)*, 2020, pp. 1–8. 1
- [18] P. Li, D. Li, W. Li, S. Gong, Y. Fu, and T. M. Hospedales, “A simple feature augmentation for domain generalization,” in *Proceedings of the IEEE/CVF international conference on computer vision*, 2021, pp. 8886–8895. 2
- [19] R. Volpi, H. Namkoong, O. Sener, J. C. Duchi, V. Murino, and S. Savarese, “Generalizing to unseen domains via adversarial data augmentation,” *Advances in neural information processing systems*, vol. 31, 2018. 2
- [20] Q. Xu, R. Zhang, Y. Zhang, Y. Wang, and Q. Tian, “A fourier-based framework for domain generalization,” in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2021, pp. 14 383–14 392. 2
- [21] X. He, J. Hu, Q. Lin, C. Luo, W. Xie, S. Song, M. H. Khan, and L. Shen, “Towards combating frequency simplicity-biased learning for domain generalization,” in *The Thirty-eighth Annual Conference on Neural Information Processing Systems*. 2
- [22] M. H. Khan, T. Zaidi, S. Khan, and F. S. Khan, “Mode-guided feature augmentation for domain generalization,” in *32nd British Machine Vision Conference, BMVC 2021*, 2021. 2
- [23] S. Hemati, G. Zhang, A. Estiri, and X. Chen, “Understanding hessian alignment for domain generalization,” in *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 2023, pp. 19004–19014. 2
- [24] Y. Wang, F. Liu, Z. Chen, Y.-C. Wu, J. Hao, G. Chen, and P.-A. Heng, “Contrastive-ace: Domain generalization through alignment of causal mechanisms,” *IEEE Transactions on Image Processing*, vol. 32, pp. 235–250, 2022. 2
- [25] D. Li, Y. Yang, Y.-Z. Song, and T. Hospedales, “Learning to generalize: Meta-learning for domain generalization,” in *Proceedings of the AAAI conference on artificial intelligence*, vol. 32, no. 1, 2018. 2
- [26] S. Chen, L. Wang, Z. Hong, and X. Yang, “Domain generalization by joint-product distribution alignment,” *Pattern Recognition*, vol. 134, p. 109086, 2023. 2
- [27] J. Cha, K. Lee, S. Park, and S. Chun, “Domain generalization by mutual-information regularization with pre-trained models,” in *European conference on computer vision*. Springer, 2022, pp. 440–457. 2
- [28] R. Yu, Y. Zhang, and J. Kwok, “Improving sharpness-aware minimization by lookahead,” in *Forty-first International Conference on Machine Learning*, 2024. 2
- [29] P. Wang, Z. Zhang, Z. Lei, and L. Zhang, “Sharpness-aware gradient matching for domain generalization,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2023, pp. 3769–3778. 2
- [30] A. Radford, J. W. Kim, C. Hallacy, A. Ramesh, G. Goh, S. Agarwal, G. Sastry, A. Askell, P. Mishkin, J. Clark *et al.*, “Learning transferable visual models from natural language supervision,” in *International conference on machine learning*. PmLR, 2021, pp. 8748–8763. 2
- [31] Y. Shu, X. Guo, J. Wu, X. Wang, J. Wang, and M. Long, “Clipood: Generalizing clip to out-of-distributions,” in *International Conference on Machine Learning*. PMLR, 2023, pp. 31 716–31 731. 2, 5
- [32] M. Wortsman, G. Ilharco, J. W. Kim, M. Li, S. Kornblith, R. Roelofs, R. G. Lopes, H. Hajishirzi, A. Farhadi, H. Namkoong *et al.*, “Robust fine-tuning of zero-shot models,” in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2022, pp. 7959–7971. 2
- [33] M. U. Khattak, H. Rasheed, M. Maaz, S. Khan, and F. S. Khan, “Maple: Multi-modal prompt learning,” in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2023, pp. 19 113–19 122. 2
- [34] S. Bose, A. Jha, E. Fini, M. Singha, E. Ricci, and B. Banerjee, “Stylip: Multi-scale style-conditioned prompt learning for clip-based domain generalization,” in *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision*, 2024, pp. 5542–5552. 2
- [35] D. Cheng, Z. Xu, X. Jiang, N. Wang, D. Li, and X. Gao, “Disentangled prompt representation for domain generalization,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024, pp. 23 595–23 604. 2
- [36] S. Addepalli, A. R. Asokan, L. Sharma, and R. V. Babu, “Leveraging vision-language models for improving domain generalization in image classification,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024, pp. 23 922–23 932. 2, 5
- [37] L. Dunlap, C. Mohri, D. Guillory, H. Zhang, T. Darrell, J. E. Gonzalez, A. Raghunathan, and A. Rohrbach, “Using language to extend to unseen domains,” *International Conference on Learning Representations (ICLR)*, 2023. 2
- [38] D. Qi, H. Zhao, A. Zhang, and S. Li, “Generalizing to unseen domains via text-guided augmentation: A training-free approach,” in *European Conference on Computer Vision*. Springer, 2024, pp. 285–300. 2, 4
- [39] V. W. Liang, Y. Zhang, Y. Kwon, S. Yeung, and J. Y. Zou, “Mind the gap: Understanding the modality gap in multi-modal contrastive representation learning,” *Advances in Neural Information Processing Systems*, vol. 35, pp. 17 612–17 625, 2022. 2, 4

- [40] J. Kim, Z. Wang, and Q. Qiu, “Constructing concept-based models to mitigate spurious correlations with minimal human effort,” in *European Conference on Computer Vision*. Springer, 2024, pp. 137–153. 3
- [41] P. Shi, M. C. Welle, M. Björkman, and D. Kragic, “Towards understanding the modality gap in clip,” in *ICLR 2023 workshop on multimodal representation learning: perks and pitfalls*, 2023. 4
- [42] S. Beery, G. Van Horn, and P. Perona, “Recognition in terra incognita,” in *Proceedings of the European conference on computer vision (ECCV)*, 2018, pp. 456–473. 4, 5
- [43] R. Gal, Y. Alaluf, Y. Atzmon, O. Patashnik, A. H. Bermano, G. Chechik, and D. Cohen-Or, “An image is worth one word: Personalizing text-to-image generation using textual inversion,” *arXiv preprint arXiv:2208.01618*, 2022. 4
- [44] Y. Zhang, N. Huang, F. Tang, H. Huang, C. Ma, W. Dong, and C. Xu, “Inversion-based style transfer with diffusion models,” in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2023, pp. 10 146–10 156. 4
- [45] R. Rombach, A. Blattmann, D. Lorenz, P. Esser, and B. Ommer, “High-resolution image synthesis with latent diffusion models,” in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2022, pp. 10 684–10 695. 4
- [46] D. Li, Y. Yang, Y.-Z. Song, and T. M. Hospedales, “Deeper, broader and artier domain generalization,” in *Proceedings of the IEEE international conference on computer vision*, 2017, pp. 5542–5550. 5
- [47] H. Venkateswara, J. Eusebio, S. Chakraborty, and S. Panchanathan, “Deep hashing network for unsupervised domain adaptation,” in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2017, pp. 5018–5027. 5, 6
- [48] K. Zhou, Y. Yang, T. Hospedales, and T. Xiang, “Deep domain-adversarial image generation for domain generalisation,” in *Proceedings of the AAAI conference on artificial intelligence*, vol. 34, no. 07, 2020, pp. 13 025–13 032. 5
- [49] A. Torralba and A. A. Efros, “Unbiased look at dataset bias,” in *CVPR 2011*. IEEE, 2011, pp. 1521–1528. 5
- [50] X. Peng, Q. Bai, X. Xia, Z. Huang, K. Saenko, and B. Wang, “Moment matching for multi-source domain adaptation,” in *Proceedings of the IEEE/CVF international conference on computer vision*, 2019, pp. 1406–1415. 5
- [51] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, and L. Fei-Fei, “Imagenet: A large-scale hierarchical image database,” in *2009 IEEE conference on computer vision and pattern recognition*. Ieee, 2009, pp. 248–255. 5
- [52] D. Hendrycks, K. Zhao, S. Basart, J. Steinhardt, and D. Song, “Natural adversarial examples,” in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2021, pp. 15 262–15 271. 5
- [53] C. Zhang, F. Pan, J. Kim, I. S. Kweon, and C. Mao, “Imagenet-d: Benchmarking neural network robustness on diffusion synthetic object,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024, pp. 21 752–21 762. 5
- [54] M. U. Khattak, S. T. Wasim, M. Naseer, S. Khan, M.-H. Yang, and F. S. Khan, “Self-regulating prompts: Foundational model adaptation without forgetting,” in *Proceedings of the IEEE/CVF international conference on computer vision*, 2023, pp. 15 190–15 200. 5
- [55] D. Lee, K. Hwang, and N. Kwak, “Unlocking the potential of unlabeled data in semi-supervised domain generalization,” in *Proceedings of the Computer Vision and Pattern Recognition Conference*, 2025, pp. 30 599–30 608. 5
- [56] G. Lee, W. Jang, J. H. Kim, J. Jung, and S. Kim, “Domain generalization using large pretrained models with mixture-of-adapters,” *arXiv preprint arXiv:2310.11031*, 2023. 5
- [57] D. Hendrycks, S. Basart, N. Mu, S. Kadavath, F. Wang, E. Dorundo, R. Desai, T. Zhu, S. Parajuli, M. Guo *et al.*, “The many faces of robustness: A critical analysis of out-of-distribution generalization,” in *Proceedings of the IEEE/CVF international conference on computer vision*, 2021, pp. 8340–8349. 5