

MoETTA: Test-Time Adaptation Under Mixed Distribution Shifts with MoE-LayerNorm

Xiao Fan^{2, 4*}, Jingyan Jiang^{1†}, Zhaoru Chen³,
Fanding Huang², Xiao Chen², Qinting Jiang², Bowen Zhang¹, Xing Tang¹, Zhi Wang²

¹School of Artificial Intelligence, Shenzhen Technology University, Shenzhen, China

²Shenzhen International Graduate School, Tsinghua University, Shenzhen, China

³College of Application Technology, Shenzhen University, Shenzhen, China

⁴College of Computer Science and Technology, Tongji University, Shanghai, China

Abstract

Test-Time adaptation (TTA) has proven effective in mitigating performance drops under single-domain distribution shifts by updating model parameters during inference. However, real-world deployments often involve mixed distribution shifts, where test samples are affected by diverse and potentially conflicting domain factors, posing significant challenges even for state-of-the-art TTA methods. A key limitation in existing approaches is their reliance on a unified adaptation path, which fails to account for the fact that optimal gradient directions can vary significantly across different domains. Moreover, current benchmarks focus only on synthetic or homogeneous shifts, failing to capture the complexity of real-world heterogeneous mixed distribution shifts. To address this, we propose **MoETTA**, a novel entropy-based TTA framework that integrates the Mixture-of-Experts (MoE) architecture. Rather than enforcing a single parameter update rule for all test samples, MoETTA introduces a set of structurally decoupled experts, enabling adaptation along diverse gradient directions. This design allows the model to better accommodate heterogeneous shifts through flexible and disentangled parameter updates. To simulate realistic deployment conditions, we introduce two new benchmarks: *potpourri* and *potpourri+*. While classical settings focus solely on synthetic corruptions (i.e., ImageNet-C), *potpourri* encompasses a broader range of domain shifts—including natural, artistic, and adversarial distortions—capturing more realistic deployment challenges. Additionally, *potpourri+* further includes source-domain samples to evaluate robustness against catastrophic forgetting. Extensive experiments across three mixed distribution shifts settings show that MoETTA consistently outperforms strong baselines, establishing new state-of-the-art performance and highlighting the benefit of modeling multiple adaptation directions via expert-level diversity.

Code — <https://github.com/AnikiFan/MoETTA>

*Work completed during an internship at Tsinghua University. Email: xiaofan140@gmail.com

†Corresponding author. Email: jiangjingyan@sztu.edu.cn
Copyright © 2026, Association for the Advancement of Artificial Intelligence (www.aaai.org). All rights reserved.

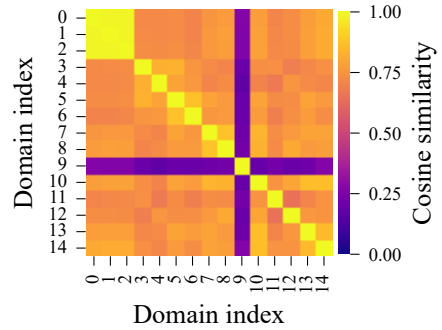


Figure 1: Cosine similarity heatmap of accumulated gradient directions $\theta_i - \theta_{\text{pre}}$, where θ_i denotes the adapted model parameters on the i -th domain of ImageNet-C (Hendrycks and Dietterich 2019) using Tent (Wang et al. 2021), and θ_{pre} is the pre-trained model parameters. Each entry at position (i, j) represents the cosine similarity between adaptation directions from domains i and j . The average cosine similarity in the lower triangle is 0.69, suggesting substantial variation across domains and highlighting the limitation of using a single adaptation direction under mixed distribution shifts.

1 Introduction

Deep learning has achieved remarkable success. However, it typically relies on the assumption that the source domain, on which the model is trained, shares the same distribution as the target domain, where the model is deployed (Ben-David et al. 2010). In practice, this assumption often fails to hold, leading to significant performance degradation under distribution shifts (Wang et al. 2024). Test-Time adaptation (TTA) (Wang et al. 2021; Niu et al. 2022; Lee et al. 2024; Niu et al. 2023; Deng et al. 2025) has emerged as a promising solution to this problem, enabling models to adapt to unseen test distributions without requiring labeled data from the target domain. TTA methods adapt the model $f(\cdot; \theta)$ by updating its parameters through the minimization of an unsupervised loss function $\mathcal{L}(\theta)$, e.g., entropy loss (Wang et al. 2021), rotation prediction loss (Gidaris, Singh, and Komodakis 2018) and contrastive-based loss (Liu et al. 2021).

Early works on TTA, such as Tent (Wang et al. 2021) and EATA (Niu et al. 2022), primarily focused on addressing single-domain distribution shifts, where all test samples originate from the same domain. Niu et al. (2023) pointed out that existing TTA methods fail under more realistic settings, with mixed distribution shifts as a representative case. For instance, in real-world deployment scenarios, inference input batches are dynamically constructed from asynchronously arriving, heterogeneous task streams, such as concurrent requests from edge devices, without guarantee of domain consistency (Cang, Chen, and Huang 2024). Although recent studies have begun to explore these challenging “in-the-wild” scenarios, a significant performance gap remains between single-domain and mixed-domain settings. A strong baseline (Deng et al. 2025) reports an accuracy drop from 71.3% to 65.9% in the latter setting, which better reflects the challenges of real-world deployment.

To understand this performance gap, we investigate the underlying causes of this limitation. As shown in Fig. 1, the average cosine similarity between the accumulated gradient directions across the 15 domains of ImageNet-C is only 0.69. This indicates that the gradient directions for different domains are misaligned. Furthermore, assuming that the domain-specific d -dimensional parameters at convergence follow a Gaussian distribution, we theoretically show that the expected cosine similarity between the corresponding accumulated gradients, from the pre-trained parameters to the converged ones, converges to $0.5 + O(1/d)$ (see App. B). This result implies the inevitability of aforementioned misalignment. More critically, in some extreme cases, e.g., domain 9 in Fig. 1, the cosine similarity with other domains can be nearly zero, leading to slow learning and poor generalization (Jacobs et al. 1991). This phenomenon reveals a fundamental limitation of current methods in mixed distribution shifts scenarios: *by enforcing a shared adaptation direction, they fail to accommodate domain-specific gradient signals, which can be inconsistent or even conflicting.*

To address this challenge, we propose **MoETTA**, a novel test-time adaptation framework that incorporates Mixture-of-Experts (MoE) paradigm (Jacobs et al. 1991; Jordan and Jacobs 1994; Shazeer et al. 2017a; Dai et al. 2024) into entropy-based test-time adaptation. Our key **insight** is that, rather than enforcing a single, unified adaptation path, leveraging a diverse set of experts to represent multiple adaptation solutions within one model is particularly advantageous for mixed distribution shifts. Specifically, MoETTA reinterprets the LayerNorm (Ba, Kiros, and Hinton 2016) parameters in Vision Transformer (ViT) (Dosovitskiy et al. 2021) as a set of structurally decoupled expert branches. These experts offer distinct parameterizations, enabling the model to accommodate diverse gradient update directions during inference. Unlike the static domain knowledge vectors used by Chen et al. (2024a), our experts are dynamically updated during inference, allowing the model to flexibly absorb distribution-specific signals. What’s more, through the routing mechanism, our methods can handle samples in a sample-wise way, which can further decompose diverse gradient directions brought by mixed distribution shifts.

To reflect more realistic deployment scenarios, we pro-

pose two new benchmarks: *potpourri* and *potpourri+*. Compared with the mixed distribution shifts setting proposed by Niu et al. (2023), which we refer to as the *classical mixed distribution shifts*, the *potpourri* setting includes samples from not only ImageNet-C (Hendrycks and Dietterich 2019) but also ImageNet-R (Hendrycks et al. 2021a), ImageNet-A (Hendrycks et al. 2021b), and ImageNet-Sketch (Wang et al. 2019), forming a heterogeneous mixture of distribution shifts. Beyond synthetic corruptions such as noise, blur, weather, and digital artifacts from ImageNet-C, *potpourri* introduces additional challenges from natural, artistic, and adversarial distortions, offering a more comprehensive testbed.

We further propose *potpourri+*, which augments *potpourri* with ImageNet validation samples to evaluate methods’ resilience to catastrophic forgetting when simultaneously handling in-distribution (ID) and out-of-distribution (OOD) data—a common challenge in practical applications. Our main contributions are summarized as follows:

- We identify a key limitation of existing TTA methods under mixed distribution shifts and address it with MoETTA, an entropy-based method that leverages expert routing to model diverse adaptation directions.
- We propose two novel evaluation settings, i.e., *potpourri* and *potpourri+*, which better simulate realistic deployment conditions involving mixed distribution shifts.
- MoETTA achieves state-of-the-art performance on both existing mixed distribution shift benchmarks and the newly proposed *potpourri* and *potpourri+* settings.

2 Related Work

Entropy-based Test-Time Adaptation. Entropy-based Test-Time Adaptation (Wang et al. 2021; Niu et al. 2022, 2023; Lee et al. 2024; Deng et al. 2025; Wang et al. 2022; Lee, Yoon, and Hwang 2024) refers to a class of TTA methods that enhance the robustness of pre-trained models under distribution shifts by minimizing prediction entropy (Grandvalet and Bengio 2004) during test-time adaptation.

Tent (Wang et al. 2021), EATA (Niu et al. 2022), SAR (Niu et al. 2023), and DeYO (Lee et al. 2024) adapt only the affine parameters of normalization layers. Among them, EATA, SAR, and DeYO improve the reliability of gradient signals through selective sample filtering, while SAR further incorporates Sharpness-Aware Minimization (SAM) (Foret et al. 2021) to enhance generalization under severe shifts. MGTTA (Deng et al. 2025) introduces a meta-gradient generator (MGG) to guide affine parameter updates.

While our method also belongs to the family of entropy-minimization-based approaches, it differs fundamentally in its architectural flexibility and update strategy. Some recent works have also explored structural modifications to enhance test-time adaptation. Notably, BECoTTA (Lee, Yoon, and Hwang 2024) proposes a MoE framework by incorporating domain-specific low-rank modules (LoRA) (Hu et al. 2022) as experts, aiming to address the challenges of Continual Test-Time Adaptation, where the dominant domain may gradually evolve over time. In contrast, our approach adopts LayerNorm as the expert unit and specifically targets the more challenging setting of mixed distribution shifts.

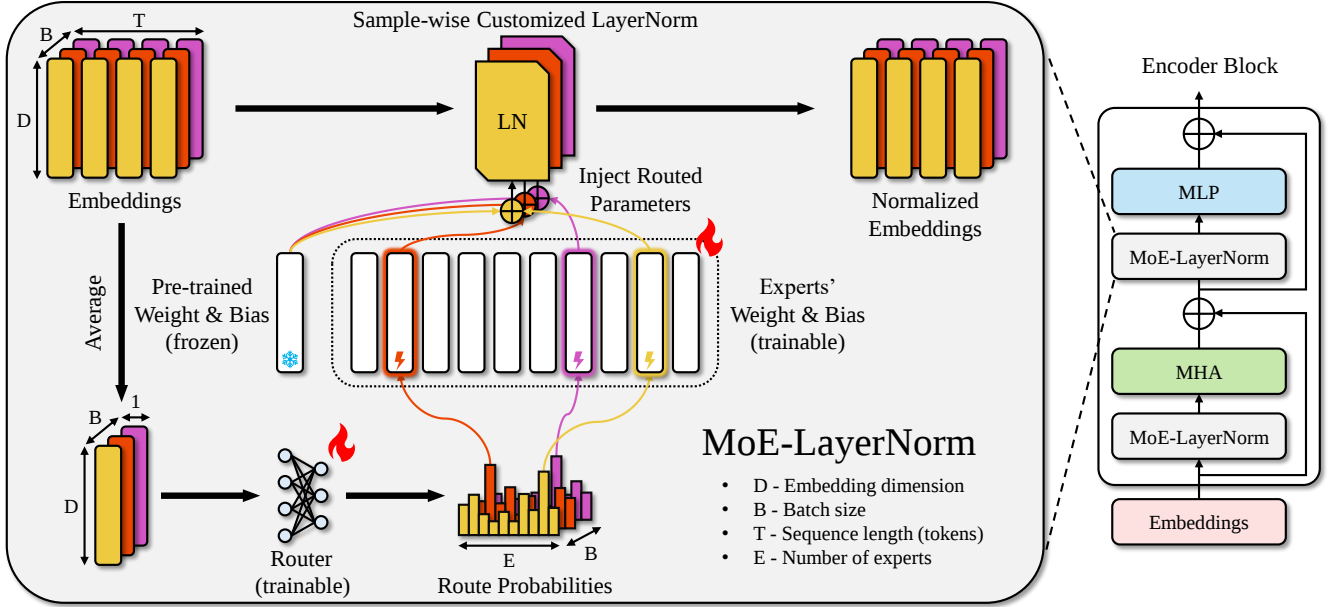


Figure 2: *Method Overview*. We replace the LayerNorm modules in the encoder blocks of a Vision Transformer (ViT) (Dosovitskiy et al. 2021) with our proposed MoE-LayerNorm. Colors of the embeddings and their routed components are used illustratively to suggest that the samples originate from different domains. For each input embedding, we first compute the mean across the token (sequence) dimension. This averaged vector is fed into a router to obtain routing probabilities, and the expert with the highest probability is selected. Its parameters are then added to the frozen pre-trained LayerNorm parameters to form a sample-specific LayerNorm. Finally, each token embedding is normalized using this customized LayerNorm. A PyTorch-style pseudo code of the forward pass for MoE-LayerNorm can be found in App. A.

Mixture of Experts. Mixture of Experts was initially proposed for its ability to handle diverse subtasks by routing inputs to multiple expert networks through a gating mechanism (Jacobs et al. 1991). In recent years, MoE has gained significant traction in natural language processing (NLP) (Shazeer et al. 2017a; Dai et al. 2024; DeepSeek-AI 2024; Fedus, Zoph, and Shazeer 2022; Jiang et al. 2024; Shazeer et al. 2017b), as it enables large-scale model capacity while maintaining a limited computational footprint during inference. For example, Fedus, Zoph, and Shazeer (2022) proposed a Mixture-of-Experts layer for natural language processing, where multiple feed-forward networks with distinct parameters are treated as experts. During inference, the top-1 experts are selected based on the output of a router, which consists of a single fully connected layer.

Similar to the Deep Mixture of Experts (DMoE) framework proposed by Eigen, Ranzato, and Sutskever (2014), our method introduces an MoE structure at every layer of the model to support diverse computation pathways. However, a key distinction lies in how the experts are implemented: while DMoE uses full linear transformations followed by nonlinearities as experts, we adopt LayerNorm instances as lightweight experts. This design choice significantly reduces the computational overhead, as only the affine parameters of LayerNorm are adapted per expert. As a result, *our approach preserves expert diversity while maintaining high efficiency, making it well-suited for test-time adaptation*.

3 Methodology

We begin this section by formally defining the mixed distribution shifts setting. We then introduce the core component of our method, the MoE-LayerNorm module, as illustrated in Fig. 2. Finally, we detail the overall adaptation procedure.

3.1 Problem Statement

Consider a pre-trained model $f(\cdot; \theta)$, where θ is learned from a labeled dataset $\{(\mathbf{x}_n, y_n)\}_{n=1}^N$ drawn from a source distribution $\mathcal{P}(\mathbf{x}, y)$. Once deployed, without access to ground-truth labels y , the model processes a stream of input sample batches $\mathcal{B}_0, \mathcal{B}_1, \dots$, each of size B , drawn from a target distribution that typically differs from $\mathcal{P}(\mathbf{x}, y)$, resulting in distribution shifts that can significantly degrade performance. In the conventional single-domain setting, all test samples—both within and across batches—are assumed to come from a single target distribution $\mathcal{Q}(\mathbf{x}, y)$, with $\mathcal{Q} \neq \mathcal{P}$. In contrast, the more realistic mixed-domain setting allows each batch to contain samples from multiple target distributions $\mathcal{Q}_1(\mathbf{x}, y), \mathcal{Q}_2(\mathbf{x}, y), \dots$, capturing diverse, domain-specific shifts encountered during deployment.

3.2 MoE-LayerNorm

Key Design Choices. Our method adapts MoE to the test-time adaptation setting through several key design choices:

1. To minimize computational overhead—a critical requirement for TTA—we designate the LayerNorm parameters

as experts instead of employing costly feed-forward networks. During adaptation, only the router and expert-specific LayerNorm parameters are updated, ensuring lightweight and efficient optimization.

2. To promote diversity across experts while maintaining efficiency, we follow the top-1 routing strategy adopted by Fedus, Zoph, and Shazeer (2022), activating only one expert per sample. This design also avoids the need to merge multiple experts' output, thereby preventing parameter interference and preserving their individuality.
3. Given the limited number of samples at test time, we construct the effective LayerNorm parameters used for normalizing each sample as the sum of the corresponding expert's parameters (initialized as zeros) selected by the router and the frozen pre-trained LayerNorm parameters, which we refer to as the *shared expert*. The inclusion of the shared expert provides domain-invariant knowledge, such as semantic representations learned from ID samples. In addition, it serves as a common foundation across all samples, which helps reduce redundancy among experts and encourages them to develop complementary adaptation behaviors (Dai et al. 2024).

Routing Mechanism. Our router is a linear projector that takes the mean sample embedding over the token dimension as input and outputs a vector whose dimension equals the number of experts. It is initialized with Xavier initialization (Glorot and Bengio 2010) at the start of adaptation.

Following Fedus, Zoph, and Shazeer (2022), in addition to task loss, we update routers with a differentiable load balancing loss. Given N experts indexed from 1 to N and the t -th batch $\mathcal{B}_t$, the auxiliary loss is computed as the scaled dot-product between two vectors $\mathbf{F}$ and $\mathbf{P}$:

$$\mathcal{L}_{\text{load balancing}} = N \times \sum_{i=1}^N \mathbf{F}_i \times \mathbf{P}_i. \quad (1)$$

Here, $\mathbf{F}_i$ denotes the fraction of samples in $\mathcal{B}_t$ routed to expert i , and $\mathbf{P}_i$ represents the average router-assigned probability for expert i w.r.t. samples in $\mathcal{B}_t$:

$$\mathbf{F}_i = \frac{1}{|\mathcal{B}_t|} \sum_{\mathbf{x} \in \mathcal{B}_t} \mathbb{I}_{\{\arg \max_k \mathbf{p}_k(\mathbf{x}) = i\}}, \mathbf{P}_i = \frac{1}{|\mathcal{B}_t|} \sum_{\mathbf{x} \in \mathcal{B}_t} \mathbf{p}_i(\mathbf{x}), \quad (2)$$

where $\mathbf{p}(\mathbf{x})$ denotes the full routing probability vector produced by the router for sample $\mathbf{x}$, while $\mathbf{p}_i(\mathbf{x})$ refers to its i -th component.

The aforementioned loss primarily encourages balanced expert utilization across samples, as it reaches its minimum value 1 when all elements in $\mathbf{F}$ are equal. A theoretical proof is provided in App. C.

However, in the test-time adaptation setting, balanced routing alone is insufficient. We also aim for the router to establish a consistent mapping that aligns inputs with appropriate experts based on shared adaptation patterns. In other words, it is desirable for samples with similar characteristics to be routed to the same expert, enabling each expert to adapt along distinct update directions rather than responding to a heterogeneous mixture of inputs.

To support such specialization, we adopt the following strategy: during the forward pass, only the expert corresponding to the maximum routing probability p is activated. To ensure that the router remains trainable under the entropy-based task loss, we apply a gradient-preserving trick by scaling the selected expert output with $p/p.\text{detach}()$. This operation does not affect the forward computation but allows gradients to flow from the entropy loss into the router. As a result, the routing decisions are optimized jointly with model predictions, encouraging experts to develop diverse and input-sensitive behaviors over time.

For the t -th batch $\mathcal{B}_t$, to maintain a balance between the load balancing loss $\mathcal{L}_{\text{load balancing}}$ and the entropy loss during adaptation, we scale $\mathcal{L}_{\text{load balancing}}$ using a dynamic weight α_t , defined as:

$$\alpha_t = \begin{cases} \lambda \times E_{\text{avg}}^0, & t = 0 \\ \alpha_{t-1} \times \frac{E_{\text{avg}}^t}{E_{\text{avg}}^{t-1}}, & t \geq 1 \end{cases}, \quad (3)$$

where λ is a hyper-parameter that controls the relative importance of the load balancing loss w.r.t. the entropy loss, and E_{avg}^t denotes the average entropy loss of all test samples up to the t -th batch.

3.3 Overall Procedure

Given an image $\mathbf{x} \in \mathbb{R}^{H \times W \times C}$, we divide this image into $N = HW/P^2$ patches $\{\mathbf{x}_p^i\}_{i=1}^N$, each with shape $P \times P \times C$. Next, we apply D filters with the same shape to each patch to create a patch embedding $\tilde{\mathbf{z}}_0^i \in \mathbb{R}^{D \times 1}$ ($i = 1, \dots, N$). We then concatenate the patch embedding with the class embedding $\mathbf{x}_{\text{class}} = \tilde{\mathbf{z}}_0^0$. Finally, we add the position encoding $\mathbf{E}_{\text{pos}} \in \mathbb{R}^{D \times (N+1)}$ to obtain the final embedding at the layer 0. With a L -layers ViT, the forward pass of MoETTA can be summarized as Eqs. (4)–(7), where MSA stands for Multi-head Self-Attention (Vaswani et al. 2017) and MLP stands for Multi-Layer Perceptron.

$$\mathbf{z}_0 = [\mathbf{x}_{\text{class}}; \mathbf{x}_p^1 \mathbf{E}; \mathbf{x}_p^2 \mathbf{E}; \dots; \mathbf{x}_p^N \mathbf{E}] + \mathbf{E}_{\text{pos}}, \quad (4)$$

$$\mathbf{z}'_\ell = \text{MSA}(\text{MoE-LayerNorm}(\mathbf{z}_{\ell-1})) + \mathbf{z}_{\ell-1}, \quad (5)$$

$$\mathbf{z}_\ell = \text{MLP}(\text{MoE-LayerNorm}(\mathbf{z}'_\ell)) + \mathbf{z}'_\ell, \quad (6)$$

$$\mathbf{p}(y|\mathbf{x}) = \text{Softmax}(\text{MLP}(\text{MoE-LayerNorm}(\mathbf{z}_L^0))), \quad (7)$$

where $\mathbf{E} \in \mathbb{R}^{D \times (P^2 C)}$, $\mathbf{E}_{\text{pos}} \in \mathbb{R}^{D \times (N+1)}$ and $\ell = 1, \dots, L$.

Inspired by Chen et al. (2024b) and Niu et al. (2022), we filter out samples with high entropy loss using a dynamic threshold and re-weight remaining samples by their entropy loss. Specifically, for the t -th batch $\mathcal{B}_t$, once the forward pass is finished, we calculate the sample-selection threshold E_{max}^t using

$$E_{\text{max}}^t = \begin{cases} E_{\text{avg}}^0, & t = 0 \\ E_{\text{max}}^{t-1} \times \frac{E_{\text{avg}}^t}{E_{\text{avg}}^{t-1}}, & t \geq 1 \end{cases}, \quad (8)$$

We then back-propagate the following multi-term loss:

$$\frac{1}{|\mathcal{S}_t|} \sum_{\mathbf{x} \in \mathcal{S}_t} \underbrace{\exp[\mathbf{E}_0 - \text{Ent}(\mathbf{x})]}_{\text{Entropy re-weighting}} \underbrace{\text{Ent}(\mathbf{x})}_{\text{Entropy loss}} + \alpha_t \sum_{i=1}^M \mathcal{L}_{\text{load balancing}}^i, \quad (9)$$

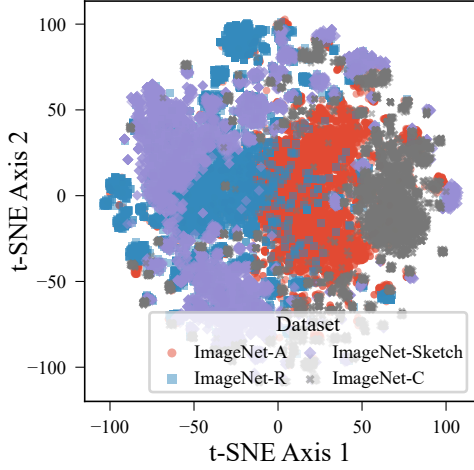


Figure 3: t-SNE (van der Maaten and Hinton 2008) projection of CLS token embeddings z_L^0 in Eq. (7), from ViT-B/16 on 7,500 samples per dataset from ImageNet-A, -R, -Sketch, and -C. While classical mixed distribution shifts (ImageNet-C only) occupy a relatively narrow region (gray), our proposed potpourri benchmark introduces greater semantic and stylistic diversity by incorporating three additional variants. This results in a more heterogeneous and realistic evaluation setting for TTA.

where $\text{Ent}(\mathbf{x})$ denotes the entropy of the posterior probability distribution $p(y|\mathbf{x})$ in Eq. (7), α_t , defined by Eq. (3), is the trade-off coefficient that balances the entropy-minimization term against the load balancing loss, M is the number of MoE-LayerNorm and $\mathcal{L}_{\text{load balancing}}^i$ is the load balancing loss for the router corresponding to the i -th MoE-LayerNorm. The set of samples that participate in the entropy term is

$$\mathcal{S}_t = \{\mathbf{x} \mid \text{Ent}(\mathbf{x}) < E_{\max}^t \wedge \mathbf{x} \in \mathcal{B}_t\}, \quad (10)$$

i.e., the “reliable” samples whose entropy falls below the current threshold. Given that there are only two degrees of freedom among the learning rate, λ in Eq. (3), and E_0 in Eq. (9), we fix E_0 as a constant and tune the others as hyperparameters throughout this work.

To clarify, among all symbols in Eqs. (1)–(10), only two need to be manually specified: λ and M , determined by the chosen MoE-LayerNorm replacement strategy. In Sec. 5.6, we further demonstrate MoETTA’s robustness to λ and provide practical guidelines for the replacement strategy. Overall, MoETTA remains flexible and easy to tune. The pseudocode of the entire procedure is presented in App. A.

4 Potpourri and Potpourri+ Benchmarks

Since Niu et al. (2023) highlighted the difficulty of addressing mixed distribution shifts and adopted a mixture of ImageNet-C as the evaluation setting, this benchmark has become the *de facto* standard for assessing the robustness of TTA methods under mixed-domain scenarios (Lee et al. 2024; Su et al. 2024; Deng et al. 2025; Tan et al. 2025). While ImageNet-C contains 15 corruption types grouped

into four categories—noise, blur, weather, and digital—it still falls short in capturing the full diversity and severity of real-world distribution shifts. For example, natural renditions can cause substantial misalignment in texture and local image statistics (Hendrycks et al. 2021a), and spurious background changes may appear even when the object of interest remains visually consistent (Wiles et al. 2022).

To close this gap, we propose the potpourri benchmark, which comprises a heterogeneous mixture of samples from ImageNet-C (Hendrycks and Dietterich 2019), ImageNet-R (Hendrycks et al. 2021a), ImageNet-Sketch (Wang et al. 2019), and ImageNet-A (Hendrycks et al. 2021b). ImageNet-R evaluates robustness against stylistic renditions, while ImageNet-A and ImageNet-Sketch probe vulnerability to spurious correlations and semantic abstraction. As shown in Fig. 3, these datasets complement ImageNet-C and together span a wider and more diverse set of domain shifts. This composition enables a more comprehensive and realistic evaluation of TTA methods.

In addition, test-time data streams may occasionally include samples from the source distribution. However, TTA models often suffer from catastrophic forgetting on such ID samples after being adapted to OOD samples (Niu et al. 2022). To address this, we introduce potpourri+, an extension of potpourri that includes validation samples from the original ImageNet dataset. It provides a more faithful assessment of a method’s ability to balance adaptation and retention during mixed-domain deployment.

5 Experiments

5.1 Experimental Protocol

We evaluate our method under three settings: the classical mixed distribution shifts, and the proposed potpourri and potpourri+ benchmarks. Details of the related datasets are provided in App. D. All evaluation settings used in this section consist of ImageNet-C at corruption severity level 5, with a batch size of 64. In this work, all pre-trained models are obtained from timm (Wightman 2019) library. Unless otherwise stated, all experiments use ViT-B/16.

We compare MoETTA with the following state-of-the-art test-time adaptation methods: Tent (Wang et al. 2021), EATA (Niu et al. 2022), SAR (Niu et al. 2023), DeYO (Lee et al. 2024), CoTTA (Wang et al. 2022), MGTTA (Deng et al. 2025), and BECoTTA (Lee, Yoon, and Hwang 2024). Additional implementation details and baseline configurations are provided in App. E.

5.2 Robustness to Mixed Distribution Shifts

We evaluate MoETTA on three mixed distribution shift benchmarks using both Transformer-based models (ViT (Dosovitskiy et al. 2021)) and convolution-based models (ConvNeXt (Liu et al. 2022)). As shown in Tab. 1, MoETTA achieves *state-of-the-art performance across all six settings*, demonstrating its effectiveness in improving the robustness of diverse architectures under mixed distribution shifts. To assess scalability, we further apply MoETTA to ViT-L/16 (304M parameters) in addition to ViT-B/16 (86M parameters); details are provided in App. I.

Model	Setting	Noadapt	Tent	EATA	CoTTA	SAR	DeYO	MGTTA	BECoTTA	Ours
ViT -B/16	Classical	55.52	63.20 _{0.08}	64.28 _{0.09}	60.53 _{0.57}	60.76 _{0.04}	63.97 _{0.04}	<u>66.20</u> _{0.01}	61.57 _{0.08}	67.20 _{0.03}
	Pot.	54.18	60.99 _{0.05}	61.99 _{0.11}	59.67 _{1.21}	58.71 _{0.03}	61.66 _{0.02}	<u>62.98</u> _{0.26}	59.08 _{0.86}	65.12 _{0.08}
	Pot.+	55.92	62.28 _{0.03}	63.17 _{0.06}	59.26 _{0.68}	59.99 _{0.07}	62.90 _{0.03}	<u>64.35</u> _{0.07}	58.87 _{3.23}	66.15 _{0.06}
Conv. -B	Classical	54.81	58.88 _{0.06}	<u>64.50</u> _{0.06}	59.65 _{0.04}	61.67 _{2.65}	64.32 _{0.03}	-	50.16 _{7.96}	67.40 _{0.02}
	Pot.	53.91	58.23 _{0.05}	<u>62.69</u> _{0.07}	58.57 _{0.00}	61.16 _{0.41}	62.46 _{0.07}	-	28.28 _{20.54}	65.70 _{0.05}
	Pot.+	55.69	59.69 _{0.04}	<u>63.94</u> _{0.07}	60.02 _{0.02}	62.72 _{0.10}	63.57 _{0.06}	-	48.92 _{9.35}	66.68 _{0.07}

Table 1: Accuracy comparison (% , $\uparrow$) under different mixed distribution settings. Conv. stands for ConvNeXt. Noadapt refers to the evaluation of the model without adaptation. Results for MGTTA applied to ConvNeXt are not reported, since pre-training MGG for architectures with multiple normalization dimensions is non-trivial. For compactness, the notation a_b denotes $a \pm b$, where b represents the standard deviation calculated from the results of three different random seeds. The best performance is highlighted in **bold** and the second best is indicated by underlining. This convention is followed in all subsequent tables.

Furthermore, *MoETTA* achieves these results without requiring any additional samples for pretraining or the calculation of statistical information. In contrast, MGTTA, a strong baseline, relies on both extra OOD and ID samples.

5.3 Computation Efficiency

Method	#Act. Params per sample	#Fwd	#Bwd	Used time
Noadapt	0	100%	0%	100%
Tent	0.04M	100%	100%	226%
EATA	0.04M	100%	80%	239%
SAR	0.03M	199%	175%	440%
DeYO	0.04M	196%	53%	317%
CoTTA	86.42M	199%	100%	798%
MGTTA	2.80M	100%	100%	227%
BECoTTA	0.13M	100%	86%	334%
Ours	0.23M	100%	76%	247%

Table 2: Computation efficiency comparison, showing per-sample activated parameter count; forward and backward propagation counts (as a percentage of total number of samples); and computation time *relative to the Noadapt baseline*. All statistics are measured under potpourri+ setting.

As shown in Tab. 2, *MoETTA* incurs modest overhead, its runtime is 247% of the model without adaptation, only slightly above the overheads of Tent, EATA, and MGTTA.

5.4 Ablation Study

Ablation Study on Loss Components in Eq. (9). Removing either the sample selection mechanism, i.e., Eq. (10), or the entropy re-weighting part in Eq. (9) leads to moderate performance drops, showing their effectiveness in filtering reliable samples and emphasizing confident predictions. In contrast, removing the load balancing loss leads to a dramatic degradation across all settings, confirming its essential role in maintaining routing stability and preventing collapse.

Ablation Study on MoE Architecture. Disabling the gradient flow from the entropy loss to the router reduces the

	Classical	Pot.	Pot.+	Avg.
Full method	67.25	65.14	66.21	66.20
Loss Components				
w/o Sample selection	<u>67.04</u>	<u>64.01</u>	57.61	<u>62.89</u>
w/o Entropy re-weight	62.86	60.51	<u>61.79</u>	61.72
w/o $\mathcal{L}_{\text{load balancing}}$	26.27	16.27	21.29	21.28
MoE Architecture				
w/o Grad to router	<u>65.17</u>	<u>62.80</u>	<u>63.92</u>	<u>63.96</u>
w/o Sample-wise router	28.69	28.60	24.96	27.42
w/o MoE-LayerNorm	22.38	17.94	26.93	22.42
w/o Layer-wise router	17.40	27.09	15.18	19.89

Table 3: Ablation study. The first group evaluates the effect of different loss components (Eq. (9)), and the second group evaluates the effect of MoE architectural choices. Accuracy (% , $\uparrow$) is reported under classical mixed distribution shifts.

model’s ability to learn informative routing strategies, resulting in a notable performance drop. Disabling sample-wise router causes the same expert to be used across all samples in a batch, which leads to severe performance collapse, highlighting the importance of input-adaptive routing. Removing layer-wise router enforces fixed expert assignment across all MoE-LayerNorm layers, limiting the model’s flexibility and hurting performance. Eliminating the entire MoE-LayerNorm module results in unstable and ineffective adaptation, confirming its necessity.

5.5 Analysis of Expert Parameter Diversity

Fig. 4 shows that, during adaptation, experts’ parameters within the same MoE-LayerNorm diverge over time. This effect is particularly pronounced in shallow layers, reflected by the decreasing cosine similarity in Fig. 4(a-b). These results suggest that experts learn distinct parameter update trajectories, rather than converging to similar representations.

Importantly, such diversity is not pre-imposed but emerges naturally from our design: experts are structurally decoupled and trained with separate gradient signals via the routing mechanism. The final-state statistics in Fig. 4(c) con-

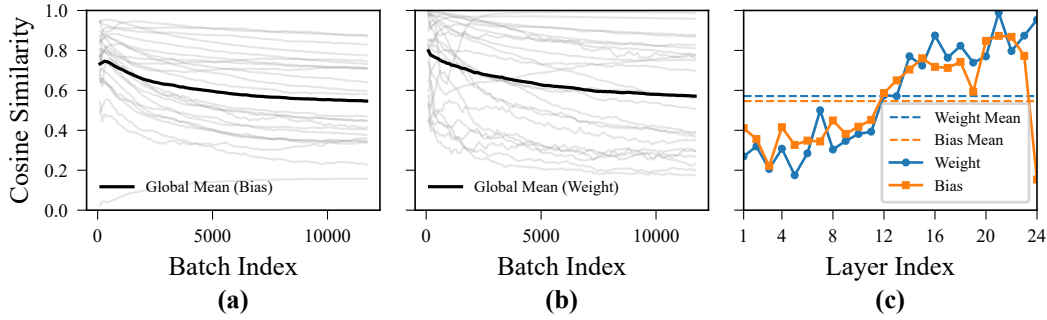


Figure 4: Evolution and final-state statistics of expert parameter similarity across MoE-LayerNorms adapted under classical mixed distribution shifts. (a–b) track the average pairwise cosine similarity of expert bias and weight parameters during adaptation, with each gray curve representing one MoE-LayerNorm and the bold black line denoting the global mean across layers. (c) reports the final average similarity for each layer at the end of adaptation phase.

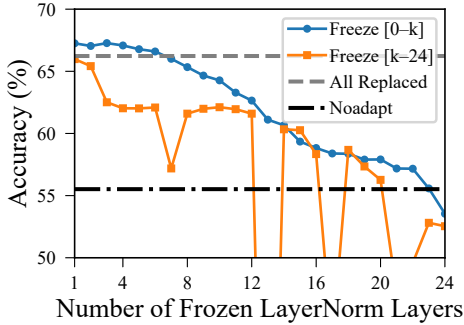


Figure 5: Effect of partially replacing LayerNorm with MoE-LayerNorm on accuracy.

firm this decoupling, with many layers exhibiting low similarity among their experts.

This parameter-level diversity allows MoETTA to internalize multiple adaptation modes within a single model, enabling it to better accommodate heterogeneous input batches without explicit domain labels. A layer-wise visualization of expert similarity is included in App. G.

5.6 Hyper-Parameter Sensitivity

In this subsection, all experiments are conducted under the classical mixed distribution shifts.

Effect of Replacement Strategy. As shown in Fig. 5, we compare two partial replacement strategies for MoE-LayerNorm: (1) freezing shallow LayerNorm layers indexed $0-k$, and (2) freezing deeper layers indexed $k-24$. Freezing the shallow layers ($k = 0, \dots, 5$) consistently outperforms full replacement. This is consistent with prior findings that shallow ViT layers primarily encode low-level, domain-invariant features such as color (Dorszewski et al. 2025), which benefit from being preserved. In contrast, allowing deeper layers—responsible for high-level semantics—to adapt leads to more effective test-time adaptation. Freezing middle layers, which are known to be critical for fine-tuning (Valeriani et al. 2023), results in a sharp performance drop. Similarly, freezing only the deeper layers hin-

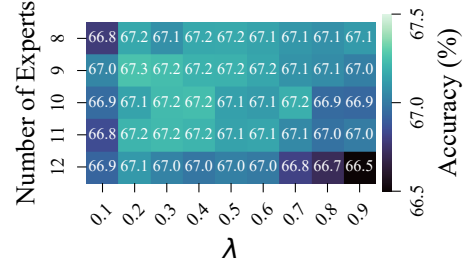


Figure 6: Grid search results over the number of experts and the trade-off coefficient λ in Eq. (3).

ders the model’s ability to adjust to distribution shifts, often leading to degraded performance or collapse.

Effect of the Number of Experts and λ in Eq. (3). As shown in Fig. 6, MoETTA performs best when the trade-off coefficient λ is small, suggesting that excessive focus on load balancing can impede effective adaptation. This underscores the value of letting the router flexibly assign experts based on input characteristics rather than enforcing uniform expert usage. The number of experts also affects performance. While MoETTA is robust across a range of choices, using 9 experts achieves the best accuracy under classical mixed distribution shifts.

6 Conclusion

We study TTA under mixed distribution shifts, where a single update direction across diverse samples often leads to suboptimal adaptation. For more realistic evaluation, we introduce two benchmarks, potpourri and potpourri+, combining multiple shift types within a single test stream.

To overcome the limitations of conventional TTA, we propose MoETTA, a lightweight framework that reparameterizes LayerNorm into a mixture of decoupled expert branches. Allowing experts to follow distinct adaptation trajectories lets MoETTA capture multiple modes of test-time behavior. Empirically, this expert diversity improves robustness and stability under complex mixed distribution shifts, consistently outperforming prior TTA methods.

Acknowledgments

This study was supported by the National Key Research and Development Program of China (Grant No. 2023YFF0905502), the National Natural Science Foundation of China (Grant Nos. 92467204 and 62472249), the Shenzhen Science and Technology Program (Grant Nos. JCYJ20220818101014030, KJZD20240903102300001, and JCYJ20250604145014018), and the Natural Science Foundation for Top Talents of Shenzhen Technology University (Grant No. GDRC202413).

References

- Ba, J. L.; Kiros, J. R.; and Hinton, G. E. 2016. Layer Normalization. *arXiv preprint arXiv:1607.06450*.
- Ben-David, S.; Blitzer, J.; Crammer, K.; Kulesza, A.; Pereira, F.; and Vaughan, J. 2010. A theory of learning from different domains. *Machine Learning*, 79: 151–175.
- Cang, Y.; Chen, M.; and Huang, K. 2024. Joint Batching and Scheduling for High-Throughput Multiuser Edge AI With Asynchronous Task Arrivals. *IEEE Transactions on Wireless Communications*, 23(10): 13782–13795.
- Chen, G.; Niu, S.; Chen, D.; Zhang, S.; Li, C.; Li, Y.; and Tan, M. 2024a. Cross-Device Collaborative Test-Time Adaptation. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*.
- Chen, Y.; Niu, S.; Xu, S.; Song, H.; Wang, Y.; and Tan, M. 2024b. Towards Robust and Efficient Cloud-Edge Elastic Model Adaptation via Selective Entropy Distillation. In *International Conference on Learning Representations*.
- Dai, D.; Deng, C.; Zhao, C.; Xu, R.; Gao, H.; Chen, D.; Li, J.; Zeng, W.; Yu, X.; Wu, Y.; Xie, Z.; Li, Y.; Huang, P.; Luo, F.; Ruan, C.; Sui, Z.; and Liang, W. 2024. DeepSeek-MoE: Towards Ultimate Expert Specialization in Mixture-of-Experts Language Models. In Ku, L.-W.; Martins, A.; and Srikumar, V., eds., *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 1280–1297. Bangkok, Thailand: Association for Computational Linguistics.
- DeepSeek-AI. 2024. DeepSeek-V3 Technical Report. *arXiv:2412.19437*.
- Deng, J.; Dong, W.; Socher, R.; Li, L.-J.; Li, K.; and Fei-Fei, L. 2009. Imagenet: A Large-Scale Hierarchical Image Database. In *CVPR*, 248–255.
- Deng, Q.; Niu, S.; Zhang, R.; Chen, Y.; Zeng, R.; Chen, J.; and Hu, X. 2025. Learning to Generate Gradients for Test-Time Adaptation via Test-Time Training Layers. In *Proceedings of the AAAI Conference on Artificial Intelligence*.
- Dorszewski, T.; Tětková, L.; Jenssen, R.; Hansen, L. K.; and Wickstrøm, K. K. 2025. From colors to classes: Emergence of concepts in vision transformers. In *World Conference on Explainable Artificial Intelligence*, 28–47. Springer.
- Dosovitskiy, A.; Beyer, L.; Kolesnikov, A.; Weissenborn, D.; Zhai, X.; Unterthiner, T.; Dehghani, M.; Minderer, M.; Heigold, G.; Gelly, S.; Uszkoreit, J.; and Houlsby, N. 2021. An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale. *ICLR*.
- Eigen, D.; Ranzato, M.; and Sutskever, I. 2014. Learning Factored Representations in a Deep Mixture of Experts. *arXiv:1312.4314*.
- Fedus, W.; Zoph, B.; and Shazeer, N. 2022. Switch transformers: scaling to trillion parameter models with simple and efficient sparsity. *J. Mach. Learn. Res.*, 23(1).
- Foret, P.; Kleiner, A.; Mobahi, H.; and Neyshabur, B. 2021. Sharpness-aware Minimization for Efficiently Improving Generalization. In *International Conference on Learning Representations*.
- Gidaris, S.; Singh, P.; and Komodakis, N. 2018. Unsupervised Representation Learning by Predicting Image Rotations. In *International Conference on Learning Representations*.
- Glorot, X.; and Bengio, Y. 2010. Understanding the difficulty of training deep feedforward neural networks. In Teh, Y. W.; and Titterton, M., eds., *Proceedings of the Thirteenth International Conference on Artificial Intelligence and Statistics*, volume 9 of *Proceedings of Machine Learning Research*, 249–256. Chia Laguna Resort, Sardinia, Italy: PMLR.
- Grandvalet, Y.; and Bengio, Y. 2004. Semi-supervised Learning by Entropy Minimization. In Saul, L.; Weiss, Y.; and Bottou, L., eds., *Advances in Neural Information Processing Systems*, volume 17. MIT Press.
- Hendrycks, D.; Basart, S.; Mu, N.; Kadavath, S.; Wang, F.; Dorundo, E.; Desai, R.; Zhu, T.; Parajuli, S.; Guo, M.; et al. 2021a. The many faces of robustness: A critical analysis of out-of-distribution generalization. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 8340–8349.
- Hendrycks, D.; and Dietterich, T. 2019. Benchmarking Neural Network Robustness to Common Corruptions and Perturbations. In *International Conference on Learning Representations*.
- Hendrycks, D.; Zhao, K.; Basart, S.; Steinhardt, J.; and Song, D. 2021b. Natural Adversarial Examples. *CVPR*.
- Hu, E. J.; Shen, Y.; Wallis, P.; Allen-Zhu, Z.; Li, Y.; Wang, S.; Wang, L.; Chen, W.; et al. 2022. Lora: Low-rank adaptation of large language models. *ICLR*, 1(2): 3.
- Jacobs, R. A.; Jordan, M. I.; Nowlan, S. J.; and Hinton, G. E. 1991. Adaptive mixtures of local experts. *Neural computation*, 3(1): 79–87.
- Jiang, A. Q.; Sablayrolles, A.; Roux, A.; Mensch, A.; Savary, B.; Bamford, C.; Chaplot, D. S.; de las Casas, D.; Hanna, E. B.; Bressand, F.; Lengyel, G.; Bour, G.; Lample, G.; Lavaud, L. R.; Saulnier, L.; Lachaux, M.-A.; Stock, P.; Subramanian, S.; Yang, S.; Antoniak, S.; Scao, T. L.; Gervet, T.; Lavril, T.; Wang, T.; Lacroix, T.; and Sayed, W. E. 2024. Mixtral of Experts. *arXiv:2401.04088*.
- Jordan, M. I.; and Jacobs, R. A. 1994. Hierarchical mixtures of experts and the EM algorithm. *Neural computation*, 6(2): 181–214.
- Lee, D.; Yoon, J.; and Hwang, S. J. 2024. BECoTTA: Input-dependent Online Blending of Experts for Continual Test-time Adaptation. In *International Conference on Machine Learning*.

- Lee, J.; Jung, D.; Lee, S.; Park, J.; Shin, J.; Hwang, U.; and Yoon, S. 2024. Entropy is not Enough for Test-Time Adaptation: From the Perspective of Disentangled Factors. In *The Twelfth International Conference on Learning Representations*.
- Liu, Y.; Kothari, P.; van Delft, B.; Bellot-Gurlet, B.; Mordan, T.; and Alahi, A. 2021. TTT++: When Does Self-Supervised Test-Time Training Fail or Thrive? In Ranzato, M.; Beygelzimer, A.; Dauphin, Y.; Liang, P.; and Vaughan, J. W., eds., *Advances in Neural Information Processing Systems*, volume 34, 21808–21820. Curran Associates, Inc.
- Liu, Z.; Mao, H.; Wu, C.-Y.; Feichtenhofer, C.; Darrell, T.; and Xie, S. 2022. A convnet for the 2020s. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 11976–11986.
- Niu, S.; Wu, J.; Zhang, Y.; Chen, Y.; Zheng, S.; Zhao, P.; and Tan, M. 2022. Efficient Test-Time Model Adaptation without Forgetting. In *The International Conference on Machine Learning*.
- Niu, S.; Wu, J.; Zhang, Y.; Wen, Z.; Chen, Y.; Zhao, P.; and Tan, M. 2023. Towards Stable Test-Time Adaptation in Dynamic Wild World. In *International Conference on Learning Representations*.
- Paszke, A.; Gross, S.; Massa, F.; Lerer, A.; Bradbury, J.; Chanan, G.; Killeen, T.; Lin, Z.; Gimelshein, N.; Antiga, L.; Desmaison, A.; Köpf, A.; Yang, E.; DeVito, Z.; Raison, M.; Tejani, A.; Chilamkurthy, S.; Steiner, B.; Fang, L.; Bai, J.; and Chintala, S. 2019. *PyTorch: an imperative style, high-performance deep learning library*, 12. Red Hook, NY, USA: Curran Associates Inc.
- Shazeer, N.; Mirhoseini, A.; Maziarz, K.; Davis, A.; Le, Q.; Hinton, G.; and Dean, J. 2017a. Outrageously large neural networks: The sparsely-gated mixture-of-experts layer. *arXiv preprint arXiv:1701.06538*.
- Shazeer, N.; Mirhoseini, A.; Maziarz, K.; Davis, A.; Le, Q.; Hinton, G.; and Dean, J. 2017b. Outrageously Large Neural Networks: The Sparsely-Gated Mixture-of-Experts Layer. In *International Conference on Learning Representations*.
- Su, Z.; Guo, J.; Yao, K.; Yang, X.; Wang, Q.; and Huang, K. 2024. Unraveling Batch Normalization for Realistic Test-Time Adaptation. In Wooldridge, M. J.; Dy, J. G.; and Natarajan, S., eds., *Thirty-Eighth AAAI Conference on Artificial Intelligence, AAAI 2024, Thirty-Sixth Conference on Innovative Applications of Artificial Intelligence, IAAI 2024, Fourteenth Symposium on Educational Advances in Artificial Intelligence, EAAI 2014, February 20-27, 2024, Vancouver, Canada*, 15136–15144. AAAI Press.
- Tan, M.; Chen, G.; Wu, J.; Zhang, Y.; Chen, Y.; Zhao, P.; and Niu, S. 2025. Uncertainty-Calibrated Test-Time Model Adaptation without Forgetting. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 1–14.
- Valeriani, L.; Doimo, D.; Cuturello, F.; Laio, A.; Ansuini, A.; and Cazzaniga, A. 2023. The geometry of hidden representations of large transformer models. In *Proceedings of the 37th International Conference on Neural Information Processing Systems, NIPS '23*. Red Hook, NY, USA: Curran Associates Inc.
- van der Maaten, L.; and Hinton, G. E. 2008. Visualizing Data using t-SNE. *Journal of Machine Learning Research*, 9: 2579–2605.
- Vaswani, A.; Shazeer, N.; Parmar, N.; Uszkoreit, J.; Jones, L.; Gomez, A. N.; Kaiser, L.; and Polosukhin, I. 2017. Attention is all you need. In *Proceedings of the 31st International Conference on Neural Information Processing Systems, NIPS'17*, 6000–6010. Red Hook, NY, USA: Curran Associates Inc. ISBN 9781510860964.
- Wang, D.; Shelhamer, E.; Liu, S.; Olshausen, B.; and Darrell, T. 2021. Tent: Fully Test-Time Adaptation by Entropy Minimization. In *International Conference on Learning Representations*.
- Wang, H.; Ge, S.; Lipton, Z.; and Xing, E. P. 2019. Learning robust global representations by penalizing local predictive power. In *Advances in Neural Information Processing Systems*, 10506–10518.
- Wang, Q.; Fink, O.; Van Gool, L.; and Dai, D. 2022. Continual Test-Time Domain Adaptation. In *Proceedings of Conference on Computer Vision and Pattern Recognition*.
- Wang, Z.; Luo, Y.; Zheng, L.; Chen, Z.; Wang, S.; and Huang, Z. 2024. In Search of Lost Online Test-Time Adaptation: A Survey. *International Journal of Computer Vision*, 133: 1106–1139.
- Wightman, R. 2019. PyTorch Image Models. <https://github.com/rwightman/pytorch-image-models>.
- Wiles, O.; Goyal, S.; Stimberg, F.; Rebuffi, S.-A.; Ktena, I.; Dvijotham, K. D.; and Cemgil, A. T. 2022. A Fine-Grained Analysis on Distribution Shift. In *International Conference on Learning Representations*.

Appendix

In this appendix, we provide additional experimental details and results. The appendix is organized into the following sections:

- Section A presents the pseudo code for the forward pass of MoE-LayerNorm, as well as the overall adaptation process of our method.
- Section B provides a proof for the convergence of the expectation of the cosine similarity of the experts.
- Section C presents a proof for the lower bound of the load balancing loss.
- Section D provides information about the datasets used in this paper.
- Section E outlines the implementation details.
- Section F describes the computing resources used throughout the experiments.
- Section G presents the visualization of the cosine similarity of the experts within each MoE-LayerNorm.
- Section H reports the results of the grid search conducted for the baselines.
- Section I reports the results using ViT-L/16 to show the scalability of MoETTA.
- Section J discusses the potential broader impacts of our research.

A Pseudocode for Adaptation Process and MoE-LayerNorm

We provide the pseudocode for the overall test-time adaptation process, along with a PyTorch (Paszke et al. 2019) style implementation of the forward pass for MoE-LayerNorm. To enhance generality, the MoE-LayerNorm pseudocode supports activating multiple experts per sample, rather than just one.

Algorithm 1: Adaptation process

Input: Test samples $\mathcal{D}_{test} = \{\{\mathbf{x}_j^{(i)}\}_{j=1}^B\}_{i=1}^N$, the pre-trained model $f(\cdot; \theta)$, number of classes N , hyper-parameters λ , E_0 .

Output: The predictions $\{\{\hat{y}_j^{(i)}\}_{j=1}^B\}_{i=1}^N$

- 1 Replace LayerNorm in $f(\cdot; \theta)$ with MoELayerNorm.
- 2 $entropys \leftarrow []$
- 3 **for** i -th batch $\{\mathbf{x}_j^{(i)}\}_{j=1}^B$ **in** $\mathcal{D}_{test}$ **do**
- 4 Calculate posterior probability : $\{\{\hat{p}_k\}_{k=1}^N\}_{j=1}^B \leftarrow f(\{\mathbf{x}_j^{(i)}\}_{j=1}^B; \theta)$
- 5 Get predictions for i -th batch : $\{\hat{y}_j^{(i)}\}_{j=1}^B \leftarrow \arg \max_k \{\{\hat{p}_k\}_{k=1}^N\}_{j=1}^B$
- 6 Calculate entropy : $\{e_j\}_{j=1}^B \leftarrow \text{entropy}(\{\{\hat{p}_k\}_{k=1}^N\}_{j=1}^B)$
- 7 $mean_entropy \leftarrow \sum_{j=1}^B e_j / B$
- 8 **if** $i==1$ **then**
- 9 $entropys.append(mean_entropy)$
- 10 $threshold \leftarrow mean_entropy$
- 11 $\alpha \leftarrow \lambda \cdot mean_entropy$
- 12 **else**
- 13 $old_avg \leftarrow mean(entropys)$
- 14 $entropys.append(mean_entropy)$
- 15 $new_avg \leftarrow mean(entropys)$
- 16 $threshold \leftarrow threshold \cdot new_avg / old_avg$
- 17 $\alpha \leftarrow \alpha \cdot new_avg / old_avg$
- 18 **end if**
- 19 $loss \leftarrow \sum_{i=1}^B [\mathbb{I}_{\{e_i < threshold\}} \exp[E_0 - e_i.detach()] \cdot e_i] / \sum_{i=1}^B \mathbb{I}_{\{e_i < threshold\}}$
- 20 $load_balancing_loss \leftarrow \text{collect_load_balancing_loss}(f(\cdot; \theta))$
- 21 $(loss + \alpha \cdot load_balancing_loss).backward()$
- 22 **end for**

Algorithm 2: MoE-LayerNorm

Input: Input $X \in \mathbb{R}^{B \times T \times D}$, expert weights $W_e \in \mathbb{R}^{E \times D}$, biases $b_e \in \mathbb{R}^{E \times D}$, shared weight $W \in \mathbb{R}^D$, bias $b \in \mathbb{R}^D$, router $g : \mathbb{R}^D \rightarrow \mathbb{R}^E$, number of activated experts k

Output: Layer-normalized tensor $Y \in \mathbb{R}^{B \times T \times D}$

```
1  $Z \leftarrow \text{mean}(X, \text{dim} = 1)$  // Shape:  $B \times D$ 
2  $P \leftarrow \text{softmax}(g(Z))$  // Router probability, shape:  $B \times E$ 
3  $\text{topk\_idx} \leftarrow \text{TopK}(P, k)$  // Shape:  $B \times k$ 
4  $\text{prob} \leftarrow \text{gather}(P, \text{topk\_idx})$  // Shape:  $B \times k$ 
5  $\text{importance} \leftarrow \text{mean}(P, \text{dim} = 0)$  // Expert mean routing prob
6  $\text{count} \leftarrow \text{bincount}(\text{topk\_idx})$  // Expert assignment count
7  $\text{load} \leftarrow \text{count} / \text{sum}(\text{count})$  // Normalized expert usage
8  $\text{aux\_loss} \leftarrow E \cdot \text{sum}(\text{importance} \cdot \text{load})$  // To be collected later
9  $c \leftarrow \text{prob} / \text{prob}.\text{detach}()$  // Forward as 1, backward to router
10  $W_{\text{sel}} \leftarrow \text{gather}(W_e, \text{topk\_idx})$  // Shape:  $B \times k \times D$ 
11  $b_{\text{sel}} \leftarrow \text{gather}(b_e, \text{topk\_idx})$  // Shape:  $B \times k \times D$ 
12  $c \leftarrow \text{unsqueeze}(c, -1)$  // Broadcast:  $B \times k \times 1$ 
13  $W_{\text{fused}} \leftarrow \text{sum}(W_{\text{sel}} \cdot c, \text{dim} = 1) + W$  // Shape:  $B \times D$ 
14  $b_{\text{fused}} \leftarrow \text{sum}(b_{\text{sel}} \cdot c, \text{dim} = 1) + b$  // Shape:  $B \times D$ 
15  $\mu \leftarrow \text{mean}(X, \text{dim} = -1, \text{keepdim} = \text{True})$  // Shape:  $B \times T \times 1$ 
16  $\sigma^2 \leftarrow \text{var}(X, \text{dim} = -1, \text{keepdim} = \text{True})$  // Shape:  $B \times T \times 1$ 
17  $X_{\text{norm}} \leftarrow (X - \mu) / \sqrt{\sigma^2 + \epsilon}$  // Shape:  $B \times T \times D$ 
18  $Y \leftarrow X_{\text{norm}} \cdot W_{\text{fused}}[:, \text{None}, :] + b_{\text{fused}}[:, \text{None}, :]$  // Shape:  $B \times T \times D$ 
```

B Theoretical Result: Cosine Similarity Expectation

We present a theoretical result on the expected cosine similarity between two independent high-dimensional Gaussian vectors, shifted by a common reference. All assumptions are explicitly stated, and a complete proof is provided below.

Proposition 1. *Let $\theta_1, \theta_2, \theta \in \mathbb{R}^d$ be independent random vectors sampled as*

$$\theta_1, \theta_2, \theta \stackrel{i.i.d.}{\sim} \mathcal{N}(\mathbf{0}, \sigma^2 I_d),$$

and define

$$\mathbf{u} := \theta_1 - \theta, \quad \mathbf{v} := \theta_2 - \theta.$$

Then, the expected cosine similarity satisfies

$$\mathbb{E}[\cos\langle \mathbf{u}, \mathbf{v} \rangle] = \frac{1}{2} + \mathcal{O}(d^{-1}), \quad \text{i.e., } \lim_{d \rightarrow \infty} \mathbb{E}[\cos\langle \mathbf{u}, \mathbf{v} \rangle] = \frac{1}{2}.$$

Moreover, this expectation is independent of σ .

Proof. **Covariance structure.** For each coordinate $i = 1, \dots, d$, we write:

$$u_i = \theta_{1,i} - \theta_i, \quad v_i = \theta_{2,i} - \theta_i.$$

By independence:

$$\text{Var}(u_i) = \text{Var}(v_i) = 2\sigma^2, \quad \text{Cov}(u_i, v_i) = \sigma^2.$$

Thus,

$$(u_i, v_i) \sim \mathcal{N}\left(\mathbf{0}, \sigma^2 \begin{pmatrix} 2 & 1 \\ 1 & 2 \end{pmatrix}\right),$$

yielding a per-coordinate correlation coefficient $\rho = \frac{1}{2}$.

High-dimensional limit. Normalize:

$$\mathbf{u}' = \frac{\mathbf{u}}{\sqrt{2}\sigma}, \quad \mathbf{v}' = \frac{\mathbf{v}}{\sqrt{2}\sigma},$$

so that all coordinates have unit variance and correlation $\rho = \frac{1}{2}$. Let

$$s_d := \frac{\mathbf{u}'^\top \mathbf{v}'}{\|\mathbf{u}'\| \|\mathbf{v}'\|} = \cos\langle \mathbf{u}', \mathbf{v}' \rangle = \cos\langle \mathbf{u}, \mathbf{v} \rangle.$$

By the strong law of large numbers:

$$\frac{\mathbf{u}'^\top \mathbf{v}'}{d} = \frac{1}{d} \sum_{i=1}^d u'_i v'_i \xrightarrow{\text{a.s.}} \rho, \quad \frac{\|\mathbf{u}'\|^2}{d} \xrightarrow{\text{a.s.}} 1, \quad \frac{\|\mathbf{v}'\|^2}{d} \xrightarrow{\text{a.s.}} 1,$$

hence

$$s_d \xrightarrow{\text{a.s.}} \rho = \frac{1}{2}.$$

Expectation. Since $s_d \in [-1, 1]$, by the dominated convergence theorem:

$$\mathbb{E}[s_d] \longrightarrow \frac{1}{2}.$$

Furthermore, a refined Berry–Esseen-type estimate gives $\mathbb{E}[s_d] = \frac{1}{2} + \mathcal{O}(d^{-1})$.

Independence of σ . We observe that in

$$\cos\langle \mathbf{u}, \mathbf{v} \rangle = \frac{\mathbf{u}^\top \mathbf{v}}{\|\mathbf{u}\| \|\mathbf{v}\|},$$

both numerator and denominator scale linearly with σ , so the expectation is invariant to the choice of σ . □

C Theoretical Result: the Lower Bound of the Load Balancing Loss

We provide a formal proof of the lower bound of the load balancing loss under Top-1 routing, following the definition introduced by Fedus, Zoph, and Shazeer (2022).

Proposition 2. *With the Top-1 routing rule (Fedus, Zoph, and Shazeer 2022), the auxiliary load balancing loss*

$$\mathcal{L}_{\text{load balancing}} = N \times \sum_{i=1}^N f_i \times P_i$$

is lower bounded by 1 for every mini-batch $\mathcal{B}$; equality holds if and only if every sample's routing probabilities are uniform:

$$p_1(\mathbf{x}) = \dots = p_N(\mathbf{x}) = \frac{1}{N}.$$

Proof. Let $\mathcal{B} = \{\mathbf{x}_1, \dots, \mathbf{x}_{|\mathcal{B}|}\}$ be a mini-batch. Denote the selected expert for each input as

$$i(\mathbf{x}) = \arg \max_k p_k(\mathbf{x}), \quad k \in \{1, \dots, N\}. \quad (1)$$

By construction, the Top-1 router guarantees that for every $\mathbf{x}$,

$$p_{i(\mathbf{x})}(\mathbf{x}) \geq \frac{1}{N}, \quad (2)$$

since the maximum of N non-negative values summing to 1 must be at least $1/N$.

Define the following quantity:

$$Z := \frac{1}{|\mathcal{B}|} \sum_{\mathbf{x} \in \mathcal{B}} p_{i(\mathbf{x})}(\mathbf{x}). \quad (3)$$

Using inequality (2), we obtain:

$$Z \geq \frac{1}{N}. \quad (4)$$

Next, rewrite Z in terms of f_i and P_i , where:

$$f_i := \frac{1}{|\mathcal{B}|} \sum_{\mathbf{x} \in \mathcal{B}} \mathbf{1}_{i(\mathbf{x})=i}, \quad P_i := \frac{1}{|\mathcal{B}|} \sum_{\mathbf{x} \in \mathcal{B}} p_i(\mathbf{x}).$$

Grouping the terms in (3) by expert index:

$$Z = \sum_{i=1}^N f_i \cdot P_i. \quad (5)$$

Combining (4) and (5), we have:

$$\sum_{i=1}^N f_i \cdot P_i \geq \frac{1}{N}. \quad (6)$$

Multiplying both sides by N , we get the desired lower bound:

$$\mathcal{L}_{\text{load balancing}} = N \cdot \sum_{i=1}^N f_i \cdot P_i \geq 1. \quad (7)$$

Tightness. Equality holds when each sample has a uniform routing probability vector:

$$p_1(\mathbf{x}) = \dots = p_N(\mathbf{x}) = \frac{1}{N}.$$

In this case, $f_i = P_i = \frac{1}{N}$ for all i , and:

$$\mathcal{L}_{\text{load balancing}} = N \cdot \sum_{i=1}^N \frac{1}{N} \cdot \frac{1}{N} = 1.$$

□

D More Details about Datasets

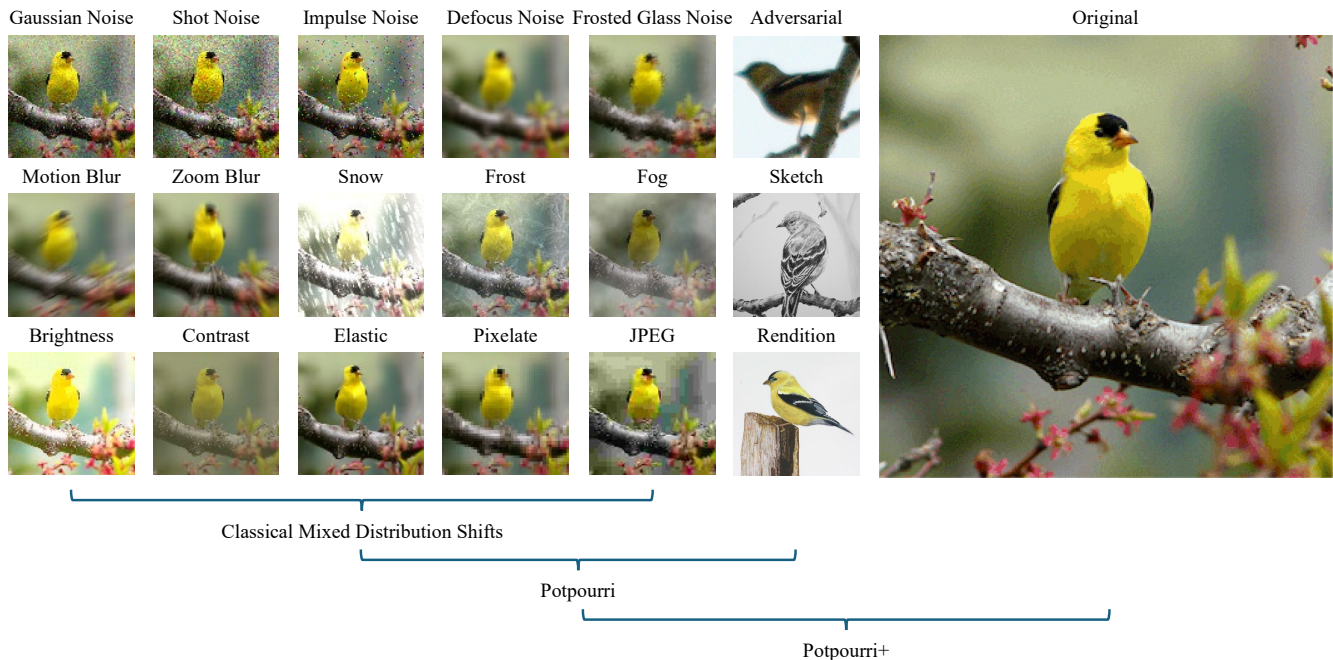


Figure 1: Overview of the corruption types and benchmark compositions used in our evaluation. The left block shows the 15 corruption types from ImageNet-C, grouped into four categories: noise, blur, weather, and digital. While classical mixed distribution shifts rely solely on these corruptions, our proposed Potpourri benchmark extends the evaluation to include samples from ImageNet-R (renditions), ImageNet-Sketch (sketches), and ImageNet-A (adversarial hard samples), thereby increasing semantic and stylistic diversity. Potpourri+ further includes clean validation images from the original ImageNet dataset to simulate real-world data streams that occasionally contain in-distribution (ID) samples.

In this paper, we utilize the mixture of following datasets to validate the robustness of our methods against mixed distribution shifts setting: ImageNet-C (Hendrycks and Dietterich 2019), ImageNet-R (Hendrycks et al. 2021a), ImageNet-Sketch (Wang et al. 2019), ImageNet-A (Hendrycks et al. 2021b) and the validation set of ImageNet (Deng et al. 2009).

Dataset	#Classes	#Samples
ImageNet (validation set)	1,000	50,000
ImageNet-C	1,000	750,000 per level
ImageNet-R	200	30,000
ImageNet-Sketch	1,000	50,889
ImageNet-A	200	7,500

Table 1: Summary of datasets used in the mixed distribution shifts evaluation.

ImageNet (Deng et al. 2009) is a large-scale image classification dataset containing over 1.2 million training images and 50,000 validation images across 1,000 categories¹. In our evaluation, we use the validation set as a representative IID dataset. Its inclusion allows us to assess whether test-time adaptation methods, when applied to out-of-distribution (OOD) data, suffer from catastrophic forgetting—i.e., a significant drop in performance on previously well-learned ID samples.

ImageNet-C (Hendrycks and Dietterich 2019) consists of 15 diverse corruption types applied to the validation images of ImageNet². These include Gaussian noise, shot noise, impulse noise, defocus blur, glass blur, motion blur, zoom blur, snow, frost, fog, brightness, contrast, elastic transform, pixelate, and JPEG compression, indexed as 0–14 in our paper. Each corruption

¹The ImageNet dataset is distributed under a custom license, available at <https://www.image-net.org/download.php>.

²The ImageNet-C dataset is distributed under CC BY 4.0 license.

type is provided at 5 severity levels, with 50,000 samples (covering 1,000 categories) per level. While the sample set remains the same across all levels, higher levels correspond to more severe distribution shifts.

ImageNet-R (Hendrycks et al. 2021a) contains 30,000 images from 200 ImageNet classes, rendered in a wide range of artistic styles such as paintings, cartoons, embroidery, and graphics³. It is designed to evaluate model robustness to style variation and domain shifts.

ImageNet-Sketch (Wang et al. 2019) consists of 50,889 black-and-white sketch images covering the same 1,000 classes as ImageNet⁴. The sketches are collected from human drawings and represent strong shifts in visual modality compared to natural images.

ImageNet-A (Hendrycks et al. 2021b) includes 7,500 natural images across 200 classes that are challenging for standard models⁵. These images are selected to be adversarial in nature—i.e., they are misclassified by high-accuracy ImageNet models while being recognizable to humans.

³The ImageNet-R dataset is distributed under MIT license.

⁴The ImageNet-Sketch dataset is distributed under MIT license.

⁵The ImageNet-A dataset is distributed under MIT license.

E More Implementation Details

All pre-trained models used for test-time adaptation in this paper are ViT-Base models obtained from the `timm` repository (Wightman 2019). To ensure reproducibility, we have released the specific model checkpoints used in our experiments.⁶ Due to the non-negligible discrepancy between the performance of the ViT-Base pre-trained weight we adopted and those of the previous studies used, as well as the sensitivity of test time adaptation methods to learning rate, we thus performed grid search on learning rate for compared methods. Detailed results of the grid search are provided in Appendix H. For other hyper-parameters, we followed prior studies when values were provided; otherwise, we adopted fair configurations—for example, by aligning the number of activated parameters across methods.

In Tables 1, we report error bars to reflect the variability across multiple runs. Specifically, each result is averaged over three independent runs with fixed random seeds: 42, 4242, and 424242. The error bars represent the *unbiased* standard deviation computed from these runs under identical settings with different random seeds.

Evaluation Metric. We use *accuracy* as the primary evaluation metric, defined as the percentage of correctly classified samples over all test samples. Formally, for N test samples, accuracy is calculated as:

$$\text{Accuracy} = \frac{1}{N} \sum_{i=1}^N \mathbb{I}[\hat{y}_i = y_i],$$

where $\hat{y}_i$ and y_i denote the predicted and ground-truth labels for the i -th sample, respectively.

Accuracy is a widely used metric in image classification tasks due to its simplicity and interpretability. It provides a direct measure of a model’s correctness, which aligns with our goal of evaluating test-time adaptation performance under mixed distribution shifts.

MoETTA(Ours). We use SGD as the update rule, with a momentum of 0.9. When batch size equals 64, we set learning rate to 1e-3. Unless otherwise specified, all the LayerNorm except for the first one in ViT are replaced by MoE-LayerNorm. For learnable parameters, inspired by Tent(Wang et al. 2021), we only update the parameters of router and experts in MoE-LayerNorm. For the classical mixed distribution shifts setting, we set λ in Eq. (3) to 0.2 and use 9 experts. For the potpourri and potpourri+ settings, we set $\lambda = 0.5$ and use 11 experts.

Tent(Wang et al. 2021). Besides learning rate, we follow all the hyper-parameters setting mentioned in Tent. Specifically, we use SGD as the update rule, and only set the affine parameters of normalization layers as trainable parameters. According to the results of the grid search, when batch size equals 64, we set the learning rate to 5e-4.

EATA(Niu et al. 2022). Besides learning rate, we follow all the hyper-parameters setting mentioned in EATA. Specifically, the entropy constant E_0 for reliable sample identification is set to $0.4 \times \ln 1000$. The ϵ for redundant sample identification is set to 0.05. The trade-off parameter β , regularization loss is set to 2000. The number of pre-collected ID test samples for Fisher importance calculation is set to 2000. The update rule is SGD, with a momentum of 0.9. Only affine parameters of normalization layers are trainable. According to the results of the grid search, when batch size equals 64, we set the learning rate to 6e-4.

CoTTA(Wang et al. 2022). Besides learning rate, we generally follow all the hyper-parameters setting mentioned in CoTTA. Specifically, we use SGD with a momentum of 0.9 as the optimizer, with an augmentation threshold p_{th} of 0.1. If images are below this threshold, we use 32 augmentations. The restoration probability is set to 0.001, and the EMA factor α for teacher update is set to 0.999. According to the results of the grid search, when batch size equals 64, we set the learning rate to 1e-3.

SAR(Niu et al. 2023). Besides learning rate, we follow all the hyper-parameters setting mentioned in SAR. Specifically, the entropy threshold, parameter ρ , moving average factor e_m and reset threshold e_0 are set to $0.4 \times \ln 1000$, 0.05, 0.9, 0.2 respectively. The update rule is SGD, with a momentum of 0.9. Only affine parameters of LayerNorm in blocks 0-8 are trainable. According to the results of the grid search, when batch size equals 64, we set the learning rate to 5e-4.

DeYO(Lee et al. 2024). Besides learning rate, we follow all the hyper-parameters setting mentioned in DeYO. Specifically, we set the parameter Ent_0 and τ_{Ent} to $0.4 \times \ln 1000$ and $0.5 \times \ln 1000$ respectively. τ_{PLPD} is set to 0.2. The update rule is SGD, with a momentum of 0.9. Only affine parameters of normalization layers are trainable. According to the results of the grid search, when batch size equals 64, we set the learning rate to 5e-5.

BECotta(Lee, Yoon, and Hwang 2024) For a fair comparison, we use BECotta without SDA initialization. MoDE is injected into every block of the ViT backbone, with six experts per block. The expert rank is set to 1, as higher values were found to cause model collapse in mixed-domain scenarios. The number of domain routers is also set to 1. Based on the results of the grid search, when the batch size equals 64, we set the learning rate to 1e-5.

⁶<https://drive.google.com/file/d/1RIrWyGnUlc12lI8BPByEmWDQI6Z3Nobk/view?usp=sharing>

MGTTA(Deng et al. 2025). Besides learning rate, we follow all the hyper-parameters setting mentioned in MGTTA. Specifically, we set the hyper-parameter λ to 0.4, the GML hidden size to 8. We adopted the pre-trained MGG released by the author.⁷ The update rule is SGD, with a momentum of 0.9. According to the results of the grid search, when batch size equals 64, we set the learning rate for ϕ^m to 3e-4.

⁷https://github.com/keikeiqi/MGTTA/blob/main/shared/mgg_ckpt.pth

F Compute Resources

All experiments were conducted on a Linux server (kernel version 5.4.0-176-generic) running Ubuntu. The machine is equipped with two Intel(R) Xeon(R) Gold 5318Y CPUs (96 cores), 251 GB of RAM, and two NVIDIA Tesla V100S GPUs with 32 GB of VRAM each. At the time of running the experiments, the NVIDIA driver version was 535.54.03, and the CUDA version was 12.2.

G Visualization of the Cosine Similarity between Different Experts

We visualize the cosine similarity of experts in different MoE-LayerNorms after adaptation under the potpourri+ setting with corruption level 5. As shown in Fig. 2 and Fig. 3, shallow MoE-LayerNorms exhibit greater variance among experts, which may be due to their role in handling texture features related to corruption.

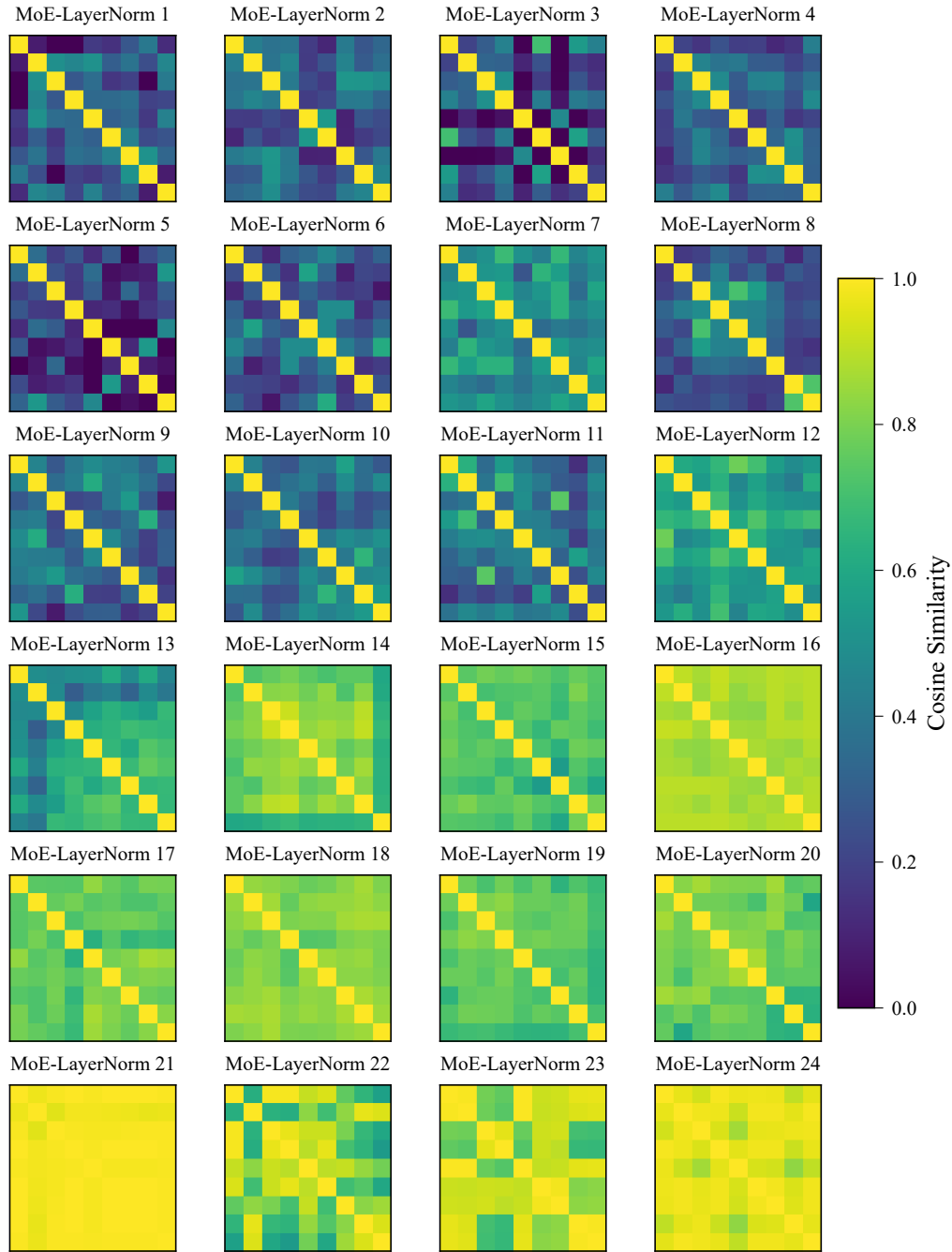


Figure 2: Cosine similarity between the expert weights within each MoE-LayerNorm layer after adaptation under the potpourri+ setting with corruption level 5. Each heatmap corresponds to one MoE-LayerNorm layer.

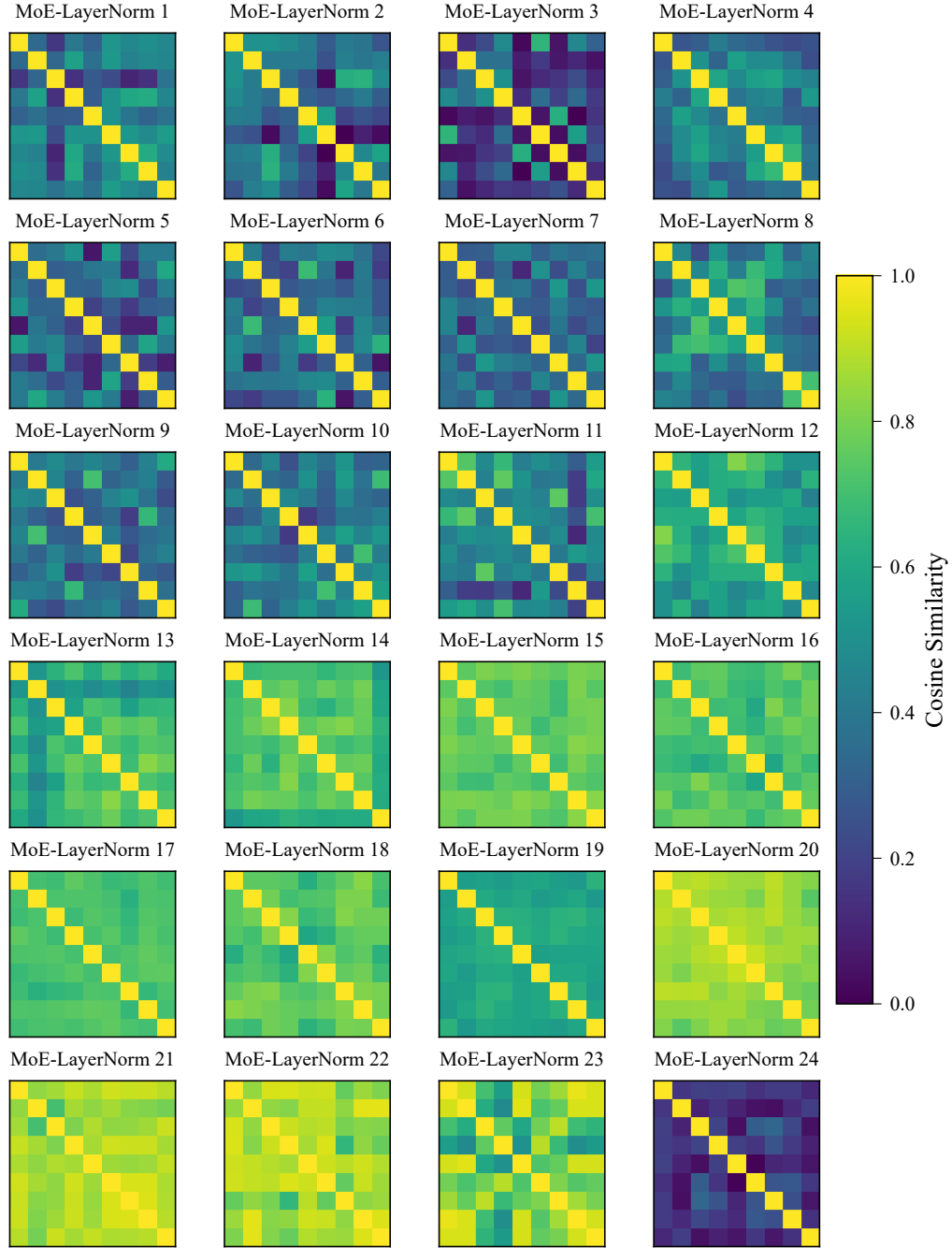


Figure 3: Cosine similarity between the expert bias within each MoE-LayerNorm layer after adaptation under the potpourri+ setting with corruption level 5. Each heatmap corresponds to one MoE-LayerNorm layer.

H Grid Search Results for Compared Methods

Since TTA methods are sensitive to the learning rate, for a fair comparison, we perform a grid search for all baselines using classical mixed distribution shifts at corruption level 5. The results are reported in Tab. 2 to Tab. 8.

Learning rate	1e-4	5e-4	1e-3
Acc.(%)	<u>61.561</u>	63.139	53.619

Table 2: Grid search results on learning rate for Tent with classical mixed distribution shifts setting of ImageNet-C(level 5).

Learning rate	5e-4	6e-4	7e-4	8e-4	9e-4	1e-3	1.1e-3	1.2e-3
Acc.(%)	62.213	64.178	<u>63.921</u>	63.892	63.569	63.407	63.238	62.639

Table 3: Grid search results on learning rate for EATA with classical mixed distribution shifts setting of ImageNet-C(level 5).

Learning rate	5e-4	1e-3	5e-3	1e-2	5e-2
Acc.(%)	<u>58.689</u>	60.273	46.530	32.058	0.243

Table 4: Grid search results on learning rate for CoTTA with classical mixed distribution shifts setting of ImageNet-C(level 5).

Learning rate	2e-4	3e-4	4e-4	5e-4	6e-4	7e-4	8e-4	9e-4	1e-3	1.1e-3	1.2e-3
Acc.(%)	60.396	60.599	60.671	60.760	<u>60.746</u>	60.712	60.725	60.712	60.637	60.598	60.676

Table 5: Grid search results on learning rate for SAR with classical mixed distribution shifts setting of ImageNet-C(level 5).

Learning rate	1e-5	5e-5	1e-4	5e-4
Acc.(%)	62.451	63.931	<u>63.111</u>	34.070

Table 6: Grid search results on learning rate for DeYO with classical mixed distribution shifts setting of ImageNet-C(level 5).

Learning rate	1e-4	2e-4	3e-4	4e-4	5e-4	6e-4	7e-4	8e-4	9e-4	1e-3	1.1e-3	1.2e-3
Acc.(%)	65.897	<u>66.195</u>	66.210	66.048	65.882	65.643	65.395	65.162	64.689	63.929	63.185	61.159

Table 7: Grid search results on learning rate for MGTТА with classical mixed distribution shifts setting of ImageNet-C (level 5).

Learning rate	1e-5	2e-5	3e-5	4e-5	5e-5	1e-4	5e-4
Acc.(%)	61.563	<u>46.187</u>	34.417	22.744	19.466	11.168	1.960

Table 8: Grid search results on learning rate for BECoTTA with classical mixed distribution shifts setting of ImageNet-C (level 5).

I Scalability

Model	Setting	Noadapt	Tent	EATA	CoTTA	SAR	DeYO	MGTTA	BECoTTA	Ours
ViT -L/16	Classical	61.0	64.7	<u>66.6</u>	-	65.5	65.6	61.7	66.2	68.8
	Pot.	59.8	62.4	64.6	-	62.0	63.6	60.6	1.5	<u>64.4</u>
	Pot.+	61.2	65.4	1.4	-	60.5	65.3	61.6	<u>66.0</u>	68.0

Table 9: Accuracy comparison (% , $\uparrow$) under different mixed distribution settings using ViT-L/16. Results for CoTTA are not reported due to out of memory error. The best performance is highlighted in **bold** and the second best is indicated by underlining. This convention is followed in all subsequent tables.

J Broader Impacts

Our work focuses on improving test-time adaptation under mixed distribution shifts, a scenario commonly encountered in real-world deployments. By enabling models to adapt dynamically to heterogeneous test samples without requiring labeled data, our method has the potential to improve the robustness and reliability of AI systems in high-stakes applications such as autonomous driving, medical diagnosis, and environmental monitoring. For example, our approach can help maintain consistent model performance across weather conditions or medical imaging devices from different institutions. These capabilities contribute toward safer and more generalizable AI deployment in practice.

We are not aware of any direct negative societal impacts associated with this work.