
Rethinking Cross-Modal Fine-Tuning: Optimizing the Interaction Between Feature Alignment and Target Fitting

Trong Khiem Tran^{1,3}, Manh Cuong Dao², Phi Le Nguyen³
Thao Nguyen Truong⁴, Trong Nghia Hoang^{1†}

¹Washington State University, ²National University of Singapore

⁴National Institute of Advanced Industrial Science and Technology

³Hanoi University of Science and Technology, [†]Corresponding author

Abstract

Adapting pre-trained models to unseen feature modalities has become increasingly important due to the growing need for cross-disciplinary knowledge integration. A key challenge here is how to align the representation of new modalities with the most relevant parts of the pre-trained model’s representation space to enable accurate knowledge transfer. This requires combining feature alignment with target fine-tuning, but uncalibrated combinations can exacerbate misalignment between the source and target feature-label structures and reduce target generalization. Existing work however lacks a theoretical understanding of this critical interaction between feature alignment and target fitting. To bridge this gap, we develop a principled framework that establishes a provable generalization bound on the target error, which explains the interaction between feature alignment and target fitting through a novel concept of feature-label distortion. This bound offers actionable insights into how this interaction should be optimized for practical algorithm design. The resulting approach achieves significantly improved performance over state-of-the-art methods across a wide range of benchmark datasets.

1 INTRODUCTION

Modern applications increasingly require transferring representational knowledge from pre-trained foundation models (FMs) to new data modalities. This allows downstream tasks to benefit from the representational knowledge embedded in the pre-trained model

when the pre-training data modalities contain relevant information. For example, in genomics, gene expression profiles can be leveraged to enrich the representation for tissue image data (Lin et al., 2024). Likewise, existing models pre-trained on broad vision (Liu et al., 2021b), language (Liu et al., 2019), or speech corpora (Radford et al., 2022) have been increasingly leveraged in domains with new data modalities not seen during pre-training. Recent studies have also shown promising transfers of vision and language models to modalities such as protein structures, cosmic ray signals, and human gestures (Shen et al., 2023).

These early results highlight the potential of cross-modal adaptation and underscore the need for more principled and generalizable approaches. Since a pre-trained FM is optimized to extract the most predictive patterns within its original representation space, positive knowledge transfer requires mapping new data modalities into the most relevant regions of that space. This raises a fundamental challenge:

How can we translate new data featuring unseen modalities into an existing pre-trained representation space to enable effective cross-modal knowledge transfer?

Addressing this challenge is non-trivial, since unlike in-modal fine-tuning, the source and target data distributions often have different statistical structures. Even when the source and target modalities are encoded into the same dimensional space, their resulting feature distributions can differ in covariance structure, higher-order interactions among covariates, and mode geometries. As the pre-trained FM implicitly leverages such distributional information to identify predictive patterns, a mismatch between the source and target representation distributions may cause it to activate spurious or irrelevant patterns, leading to negative transfer. Aligning the representation of new data modalities with relevant distributional geometries of the pre-trained representation space is therefore critical to ensure positive transfer (see Fig. 2).

On the other hand, the pre-trained representation space often contains a broad landscape of heterogeneous geometries centered around different modes, many of which may be irrelevant or poorly aligned with the target task. Thus, to avoid negative transfer that harms performance, feature alignment must also be guided by target fitting using fine-tuning data. This guided alignment is non-trivial to formalize, as the interaction between feature alignment and target fitting and its effect on target generalization is not well understood. Existing work addressing this challenge has largely focused on heuristic combinations of feature alignment and target fitting, without providing guarantees on generalization performance for the (downstream) target task (Shen et al., 2023; Cai et al., 2024; Ma et al., 2024).

Notably, ORCA (Shen et al., 2023) aligns source and target representation distributions using optimal transport before fine-tuning the entire network, while PARE (Cai et al., 2024) introduces a gating mechanism to combine source and target features during fine-tuning. These methods have demonstrated promising empirical results on benchmarks such as NasBench360 (Tu et al., 2022), but depend on heuristic formulations without explicitly modeling the interaction between feature alignment and target fitting. MoNA (Ma et al., 2024) characterizes this interaction through a bi-level optimization framework. The inner loop identifies the target-optimal predictor for a given feature embedder, while the outer loop updates the embedder so that its combined representations and predictions align with the source task’s feature-label semantics. This approach aims to reduce the misalignment between the source and target feature-label semantics under a candidate target representation using heuristic alignment measures, while fitting to target data. However, despite its empirical gains, the reliance on heuristic metrics and bi-level design leaves the theoretical link to optimal generalization on the target task unexamined. It remains unclear whether this characterization captures the most effective interaction between feature alignment and target fitting for cross-modal knowledge transfer.

To bridge this gap, we introduce a principled framework that establishes a provable bound on the generalized target error, capturing the interaction between feature alignment and target fitting via a **feature-label distortion** concept. This distortion quantifies the complexity of the probabilistic transport map between the source and target feature-label predictive distributions under a given target feature representation. Intuitively, it measures the cross-modal transferability under a given target representation. A large distortion means low transferability which will cause

target fitting to overfit when fine-tuning data is limited, thus decreasing generalized performance. This provides actionable insights into how this interaction should be optimized, offering a theoretically grounded guide for algorithm design. The above is substantiated with the following technical contributions:

Theoretical Analysis. We develop a theoretical bound that decomposes the generalized target error into: (i) the source task error, which serves as a fixed overhead; (ii) a distributional distance between the source and target representation distributions (i.e., **feature alignment**); (iii) the minimum entropy over a space of probabilistic transport maps between the source and target feature-label conditional distributions (i.e., **feature-label distortion**); and (iv) an alignment term that reflects how well the target predictor follows this transport (i.e., **target fitting**). To the best of our knowledge, this is the first generalization bound that captures the influence of both the pre-trained model’s quality and the interaction between feature alignment and target fitting on cross-modal fine-tuning performance (Section 2).

Algorithm Design. We develop a practical algorithm to address the intractability of optimizing the (ii) feature alignment and (iii) feature-label distortion terms over the space of target feature representations and the induced probabilistic transport plans between the source and target feature-label conditional distributions while performing (iv) target fitting. Our approach constructs an optimizable surrogate that serves as a medium for selectively transferring source knowledge relevant to the target task. This surrogate is optimized in a preparatory stage to guide the initialization of the main fine-tuning procedure, enabling efficient and targeted cross-modal adaptation (Section 3).

Evaluation. We evaluate our approach on two comprehensive cross-modal fine-tuning benchmarks: (1) NAS-Bench-360 (Tu et al., 2022), which covers a broad set of tasks across ten distinct data modalities; and (2) PDEBench (Takamoto et al., 2022), which assesses model adaptation to simulated data derived from diverse families of partial differential equations (PDEs). Across both benchmarks, our method consistently outperforms recent state-of-the-art baselines, including ORCA (Shen et al., 2023), PARE (Cai et al., 2024), and MoNA (Ma et al., 2024), on a significant majority of tasks. These results underscore the importance of designing algorithms guided by an explicit generalization bound framework (Section 4).

2 THEORETICAL ANALYSIS

This section presents our main theoretical result. We begin by formalizing the cross-modal fine-tuning problem setting and introducing the key notations. We then establish a generalization bound that characterizes how the interaction between feature alignment and target fitting, under a given target feature representation, affects the generalized target performance. This provides a principled foundation for understanding and optimizing cross-modal adaptation (Section 3).

2.1 Problem Setting and Notations

Let $M_s \triangleq (\theta, p_s(z \mid \theta(\mathbf{x})))$ denote the learned embedder θ and the prediction map $p_s(z \mid \theta(\mathbf{x}))$ of an FM pre-trained on a source dataset $(\mathbf{X}, \mathbf{z}) = \{(\mathbf{x}_i, z_i)\}_{i=1}^n \sim D_s(\mathbf{x}, z)$. During fine-tuning, $(\mathbf{X}, \mathbf{z})$ might not be accessible, but we assume access to an in-modal proxy dataset $(\mathbf{X}^s, \mathbf{z}^s) = \{\mathbf{x}_i^s, z_i^s\}_{i=1}^m \sim D_s(\mathbf{x}, z)$ which is sampled from the same distribution.

Let $(\mathbf{X}^\tau, \mathbf{z}^\tau) = \{\mathbf{x}_i^\tau, z_i^\tau\}_{i=1}^k \sim D_\tau(\mathbf{x}', z')$ denote the target fine-tuning dataset sampled from another data distribution D_τ over unseen modalities $\mathbf{x}'$ and label z' . We want to construct a target model $M_\tau \triangleq (\phi, p_\tau(z' \mid \phi(\mathbf{x}')))$ based on M_s and $(\mathbf{X}^\tau, \mathbf{z}^\tau) \sim D_\tau(\mathbf{x}', z')$ in a principled manner.

To achieve this, we will establish a mathematical connection (see Theorem 3) between the generalized source/target losses (see Definition 1) and the alignment of their feature distributions (see Definition 2) under feature maps θ and ϕ .

Definition 1 (Generalized Error) *The generalized source/target errors under feature maps θ/ϕ are:*

$$\text{err}_s(\theta) \triangleq -\mathbb{E}_{(\mathbf{x}, z) \sim D_s} \left[\log p_s(z \mid \theta(\mathbf{x})) \right], \quad (1)$$

$$\text{err}_\tau(\phi) \triangleq -\mathbb{E}_{(\mathbf{x}', z') \sim D_\tau} \left[\log p_\tau(z' \mid \phi(\mathbf{x}')) \right], \quad (2)$$

which are the expected source and target prediction losses at $(\mathbf{x}, z) \sim D_s$ and $(\mathbf{x}', z') \sim D_\tau$.

Definition 2 (Feature Distribution) *The feature distributions $D_s^\theta(\mathbf{u})/D_\tau^\phi(\mathbf{u})$ are defined as the push-forward of the source's and target's marginal input distributions $D_s(\mathbf{x})/D_\tau(\mathbf{x}')$ under the source and target feature maps, $\mathbf{u} = \theta(\mathbf{x})$ and $\mathbf{u} = \phi(\mathbf{x}')$, respectively.*

We also use $D_s^\theta(z \mid \mathbf{u})$ and $D_\tau^\phi(z' \mid \mathbf{u})$ to denote the source's and target's feature-label conditionals induced from the data distributions $D_s(\mathbf{x}, z)$ and $D_\tau(\mathbf{x}', z')$ under the source feature map $\mathbf{u} = \theta(\mathbf{x})$ and the target feature map $\mathbf{u} = \phi(\mathbf{x}')$, respectively.

2.2 Main Result

Our main result characterizes the generalized target loss $\text{err}_\tau(\phi)$ in terms of the following key quantities:

Overhead. The generalized source loss $\text{err}_s(\theta)$.

Feature Alignment (FA). A function of the distributional distance between source and target distributions, $D_s^\theta(\mathbf{u})$ and $D_\tau^\phi(\mathbf{u})$, over feature map θ and ϕ . (see Definition 4).

Feature Label Distortion (FLD). A minimum entropy of a valid transport plan $\Lambda_{\mathbf{u}}^*(z' \mid z)$ over label pairs (z', z) that maps the source conditional $D_s^\theta(z \mid \mathbf{u})$ to the target conditional $D_\tau^\phi(z' \mid \mathbf{u})$ over the representation $\mathbf{u} = \phi(\mathbf{x}')$ produced by the target feature map ϕ (see Definition 5).

Target Fitting (FT). A prediction alignment between the target predictor $p_\tau(z' \mid \mathbf{u} = \phi(\mathbf{x}'))$ and oracle predictor $D_\tau^\phi(z' \mid \mathbf{u} = \phi(\mathbf{x}'))$ under feature map ϕ over the presentation $\mathbf{u} = \phi(\mathbf{x}')$ (see Definition 6).

An informal statement of our result is stated below.

Theorem 3 (Informal Statement) *Under arbitrary feature map θ and ϕ , we have:*

$$\text{err}_\tau(\phi) \leq \text{err}_s(\theta) + \text{Feature-Label Distortion} \quad (3) \\ + \text{Feature Alignment} + \text{Target Fitting}.$$

This result reveals how the pre-trained model's quality and the interplay between feature alignment and target fitting shape the cross-modal fine-tuning performance via a feature-label distortion measurement (Definition 5). The source loss acts as a fixed overhead, reflecting the influence of the source model. For FMs, this loss is often negligible. The remaining terms suggest that aligning features without careful calibration may inadvertently increase the semantic gap between source and target feature-label structures. For instance, when alignment induces representations that enlarge this gap, it can harm generalization by steering target fitting toward overfitting the prediction map to the target data in order to compensate for the poorly aligned representation structure.

This insight is made precise via the below formal definitions and theorem statement (see Theorem 7).

Definition 4 (Feature Alignment) *Let Δ denote the set of cost metrics δ (on the pre-trained representation space) such that the cross-entropy of the source prediction $\ell_s(\mathbf{u}) \triangleq -\mathbb{E}_{D_s^\theta(z \mid \mathbf{u})} \log p_s(z \mid \mathbf{u})$ is τ_δ -Lipschitz with δ :*

$$|\ell_s(\mathbf{u}_1) - \ell_s(\mathbf{u}_2)| \leq \tau_\delta \cdot \delta(\mathbf{u}_1, \mathbf{u}_2). \quad (4)$$

The feature alignment under target feature map ϕ is

$$\mathbf{FA}(\phi, \theta) \triangleq \min_{\delta \in \Delta} \left\{ \tau_\delta \cdot W_\delta \left(D_\tau^\phi(\mathbf{u}), D_s^\theta(\mathbf{u}) \right) \right\}. \quad (5)$$

where W_δ is Wasserstein-1 distance with cost metric δ (see definition in Appendix D.1).

Definition 5 (Feature-Label Distortion) Let C_u^* denote the set of valid transport plans that satisfy

$$D_\tau^\phi(z' | \mathbf{u}) = \mathbb{E}_z \left[\Lambda_u^*(z' | z) \right] \text{ with } z \sim D_s^\theta(z | \mathbf{u}). \quad (6)$$

The feature-label distortion at representation $\mathbf{u}$ is

$$\mathbf{FLD}(\mathbf{u}) \triangleq \min_{\Lambda_u^* \in C_u^*} \mathbb{E}_z \left[\mathbb{H} \left[\Lambda_u^*(z' | z) \right] \right]. \quad (7)$$

with $z \sim D_s^\theta(z | \mathbf{u})$. This characterizes the transferability of source knowledge θ to target task ϕ at $\mathbf{u}$.

Definition 6 (Target Fitting) Let C_u denote the set of valid transport plans $\Lambda_u(z' | z)$ that satisfy:

$$p_\tau(z' | \mathbf{u}) = \mathbb{E}_z \left[\Lambda_u(z' | z) \right] \text{ with } z \sim D_s^\theta(z | \mathbf{u}). \quad (8)$$

The alignment of the target prediction map $p_\tau(z' | \mathbf{u})$ with the oracle predictor $D_\tau^\phi(z' | \mathbf{u})$ at $\mathbf{u} = \phi(\mathbf{x}')$ is

$$\mathbf{TF}(\mathbf{u}) \triangleq \min_{\Lambda_u \in C_u} \mathbb{E}_z \left[\mathbb{KL} \left(\Lambda_u^+(\cdot | z) \| \Lambda_u(\cdot | z) \right) \right] \text{ where} \\ \Lambda_u^+(\cdot | z) = \arg \min_{\Lambda_u^* \in C_u^*} \mathbb{E}_z \left[\mathbb{H} \left[\Lambda_u^*(z' | z) \right] \right]. \quad (9)$$

with $z \sim D_s^\theta(z | \mathbf{u})$. Theorem 3 is formally stated as:

Theorem 7 (Formal Statement) Given the above technical specification of feature alignment, feature-label distortion, and target fitting, we have

$$\text{err}_\tau(\phi) \leq \text{err}_s(\theta) + \mathbf{FA}(\phi, \theta) \\ + \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \left[\mathbf{FLD}(\mathbf{u}) + \mathbf{TF}(\mathbf{u}) \right]. \quad (10)$$

A detailed proof is provided in Appendix A. Our empirical inspection in Fig. 5, Appendix B further shows that the bound is sufficiently tight.

Theorem 7 operationalizes the earlier intuition by providing a complete algorithmic structure to measure the semantic gap in cross-modal fine-tuning under a candidate target representation, captured through both feature alignment and feature-label distortion. While prior work has relied on distributional alignment to support knowledge transfer into target fitting, the role of feature-label distortion in shaping transferability has been overlooked. This term exposes the representational discrepancy between the feature-label structures

of the source and target domains. Lower feature-label distortion suggests that the target label can be more readily inferred from the source label information, establishing a consistent pattern that improves transferability from source to target. This insight reveals that minimizing feature alignment alone may not be sufficient for effective transfer. We will build on this to design an algorithm that incorporates source-informed regularization to steer away from representations that induce large semantic gap, guiding fine-tuning toward more transferable solutions (Section 3).

3 ALGORITHM DESIGN

This section introduces a new cross-modal fine-tuning algorithm named **RECRAFT** – **RE**thinking **C**ross-**Mod**al **F**ine-**T**uning – which is designed to bridge the semantic gap between source and target tasks in a cross-modal context by optimizing the interaction between feature alignment and target fitting. It is guided by the result of Theorem 7, which bounds the generalized target error by a combination of feature alignment (**FA**) in Eq. (5), target fitting (**TF**) in Eq. (9), as well as their interaction measured via feature-label distortion (**FLD**) in Eq. (7). RECRAFT optimizes this bound via a two-stage workflow as illustrated in Figure 1 below. Stage 1 learns the optimal target feature map ϕ via minimizing a combination of **FA** and **FLD** which defines the source-target semantic gap under ϕ (Section 3.1). Stage 2 optimizes a target prediction map to minimize the target fitting term **TF** based on the learned feature map ϕ (Section 3.2).

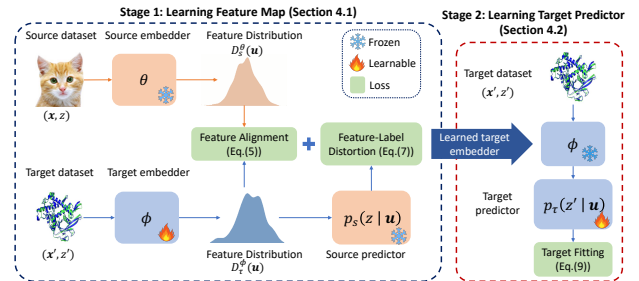


Figure 1: Overview of the RECRAFT algorithm.

To elaborate on this design, we note that a direct minimization of the theoretical bound in Eq. (10) of Theorem 7 is however unstable due to the entangled effect of optimizing both the target prediction map $p_\tau(z' | \mathbf{u})$ and feature map $\mathbf{u} = \phi(\mathbf{x}')$ on the oracle and learnable transport sets C_u^* (see Definition 5) and C_u (see Definition 6) in complex and interdependent ways. In particular, changes in the representation ϕ simultaneously alter the alignment between source and target feature distributions and reshape the transport land-

scape $C_{\mathbf{u}}^*$ on which the target predictor $p_{\tau}(z' | \mathbf{u})$ is optimized via minimizing $\mathbf{TF}$ in Eq. (9).

This coupling complicates optimization as it continually moves the optimization target for $p_{\tau}(z' | \mathbf{u})$ which destabilizes convergence. To mitigate this, we adopt a two-stage approach that decomposes the bound minimization into (1) finding a feature map ϕ that minimizes the semantic gap $\mathbf{FA}(\phi, \theta) + \mathbb{E}_{D_{\tau}^{\phi}(\mathbf{u})}[\mathbf{FLD}(\mathbf{u})]$ between the source and target task, thus maximizing transferability (Section 3.1); and (2) learning a predictor $p_{\tau}(z' | \mathbf{u})$ based on the learned ϕ (Section 3.2) via minimizing $\mathbb{E}_{D_{\tau}^{\phi}(\mathbf{u})}[\mathbf{TF}(\mathbf{u})]$. This decomposition stabilizes the optimization by removing interdependencies: the first stage depends only on ϕ while the second stage optimizes p_{τ} given ϕ , avoiding a moving target. Figure 1 provides an overview of our algorithm. Its detailed pseudocode is provided in Appendix G.

3.1 Learning Feature Map

To minimize the semantic gap between source and target, we aim to solve the following:

$$\phi = \operatorname{argmin}_{\phi} \left\{ \mathbf{FA}(\phi, \theta) + \mathbb{E}_{D_{\tau}^{\phi}(\mathbf{u})}[\mathbf{FLD}(\mathbf{u})] \right\}. \quad (11)$$

As mentioned earlier, this reveals a more principled approach to conditioning the feature map. Prior work such as (Shen et al., 2023) often minimizes $\mathbf{FA}$ primarily to align the source and target representations, without accounting for its effect on feature-label distortion ($\mathbf{FLD}$), which captures the semantic misalignment between the source and target feature-label structures induced by ϕ . As $\mathbf{FLD}$ is incorporated into Eq. (11), our formulation discourages feature maps that increase semantic misalignment which would cause the target fitting stage to compensate improperly during training and reduce transfer performance. To minimize Eq. (11), we develop effective optimization surrogates for both $\mathbf{FA}$ and $\mathbf{FLD}$ as they are not directly tractable. This is detailed next.

A. Feature Alignment Loss. Following Definition 4,

$$\mathbf{FA}(\theta, \phi) = \min_{\delta \in \Delta} \left\{ \tau_{\delta} W_{\delta}(D_{\tau}^{\phi}(\mathbf{u}), D_s^{\theta}(\mathbf{u})) \right\}, \quad (12)$$

where Δ is the set of cost metrics δ for the Wasserstein distance W_{δ} such that the cross-entropy source prediction $\ell_s(\mathbf{u}) \triangleq -\mathbb{E}_{D_s^{\theta}(z|\mathbf{u})}[\log p_s(z | \mathbf{u})]$ is τ_{δ} -Lipschitz with respect to $\delta(\mathbf{u}, \mathbf{u}')$ (see Definition 4).

To effectively constrain the search over Δ to metric regimes that impose low τ_{δ} on $\ell_s(\mathbf{u})$, our main approach is to view $\tau_{\delta} = \omega$ as a hyperparameter to be determined using a proxy dataset $P_s = (\mathbf{X}^s, \mathbf{z}^s)$ of the source task. Given ω , we will adapt the source

prediction map $p_s(z | \mathbf{u}) \simeq p_s(z | \mathbf{u}; \gamma)$ so that its imposed Lipschitz constant on $\ell_s(\mathbf{u})$ under feature map $\mathbf{u} = \theta(\mathbf{x})$ is around $O(\omega)$. This is achieved via

$$\operatorname{minimize}_{\mathbf{x} \sim P_s} \mathbb{E} \max \left(0, \|\nabla_{\mathbf{u}} \ell_s(\theta(\mathbf{x}); \gamma)\|_{\delta} - \omega \right)^2, \quad (13)$$

with respect to δ and ω . Here, γ can be selected as the parameters of the last layer in the pre-trained prediction map $p_s(z | \mathbf{u})$. This generalizes a prior practice on Lipschitz conditioning in (Shen et al., 2018). The hyperparameter ω can be selected to have smallest value without degrading the predictive performance of the source model on the proxy dataset P_s . In our experiment, this optimal value ranges between 0.3 and 0.5 across different source models under the Euclidean metric $\delta \equiv \ell_2$ (see Appendix F). Thus, the feature alignment ($\mathbf{FA}$) can be surrogated with:

$$\mathbf{FA}(\theta, \phi) \simeq L_{\mathbf{FA}}(\phi) \triangleq \omega \cdot W_{\ell_2}(D_{\tau}^{\phi}(\mathbf{u}), D_s^{\theta}(\mathbf{u})), \quad (14)$$

where W_{ℓ_2} is the Wasserstein-1 distance with norm-2 cost metric ℓ_2 . See Appendix D.2 for more details.

B. Feature-Label Distortion Loss. We will now construct a surrogate for feature-label distortion,

$$\mathbb{E}_{D_{\tau}^{\phi}(\mathbf{u})}[\mathbf{FLD}(\mathbf{u})] = \mathbb{E}_{(\mathbf{u})} \left[\min_{\Lambda_{\mathbf{u}}^*} \mathbb{E}_z \left[\mathbb{H}[\Lambda_{\mathbf{u}}^*(\cdot | z)] \right] \right] \quad (15)$$

$$\leq \mathbb{H}_{(\phi, p_s)}[Z' | Z, \mathbf{U}] \quad (16)$$

$$\leq \mathbb{H}_{(\phi, p_s)}[Z' | Z] \quad (17)$$

$$= \mathbb{H}_{(\phi, p_s)}[Z', Z] - \mathbb{H}_{(\phi, p_s)}[Z], \quad (18)$$

where the first inequality hold due to the minimum over the choice of the oracle transport $\Lambda_{\mathbf{u}}^*(z' | z)$ and the definition of conditional entropy. The second inequality holds due the information-never-hurt property of entropy (Cover and Thomas, 2006). The last equality holds from the chain rule. Eq. (18) allows us to bypass the inaccessible oracle transport $\Lambda_{\mathbf{u}}(z' | z)$ and estimate it using empirical methods (Nguyen et al., 2020). For each data point $(\mathbf{x}', z')$ in the target dataset, we can generate a pseudo source label for it via $z \sim P_s(z | \mathbf{u})$ under the target feature map $\mathbf{u} = \phi(\mathbf{x}')$. Using these statistics, we can approximate $P_{(\phi, p_s)}(z, z') \simeq C(z, z')/\kappa$ where $C(z, z')$ counts the number of times we observe (z, z') via the above pseudo source simulation; and κ is the number of target data points. Likewise, we estimate $P_{(\phi, p_s)}(z) \simeq \sum_{z'} P_{(\phi, p_s)}(z, z')$. This allows us to approximate

$$\begin{aligned} \mathbb{H}_{(\phi, p_s)}[Z' | Z] &= - \sum_z \sum_{z'} P_{(\phi, p_s)}(z, z') \log P_{(\phi, p_s)}(z, z') \\ \mathbb{H}_{(\phi, p_s)}[Z] &= - \sum_z P_{(\phi, p_s)}(z) \log P_{(\phi, p_s)}(z). \end{aligned} \quad (19)$$

The surrogate for **FLD** is thus defined as

$$\begin{aligned} \mathbb{E}_{D_{\tau}^{\phi}(\mathbf{u})}[\mathbf{FLD}(\mathbf{u})] &\simeq L_{\mathbf{FLD}}(\phi) \triangleq \mathbb{H}_{(\phi, p_s)}[Z' | Z] \quad (20) \\ &= \mathbb{H}_{(\phi, p_s)}[Z', Z] - \mathbb{H}_{(\phi, p_s)}[Z]. \quad (21) \end{aligned}$$

Combining Eq. (14) and Eq. (21), we obtain the following surrogate for Eq. (11),

$$\phi = \operatorname{argmin}_{\phi} \left(L_{\mathbf{FA}}(\phi) + L_{\mathbf{FLD}}(\phi) \right), \quad (22)$$

which can be minimized effectively using standard numerical optimization method.

3.2 Learning Target Predictor

Given the learned target feature map ϕ (Section 3.1), we parameterize the target predictor,

$$p_{\tau}(z' | \mathbf{u}) = \mathbb{E}_z[\Lambda_{\mathbf{u}}(z' | z)] \text{ with } z \sim D_s^{\theta}(z | \mathbf{u}), \quad (23)$$

and a learnable transport $\Lambda_{\mathbf{u}}(z' | z)$. For convenience, we can also approximate $D_s^{\theta}(z | \mathbf{u})$ with $p_s(z | \mathbf{u})$ since the predictive map p_s of a pre-trained foundation model often capture well the source’s feature-label conditional. Consequently, the learning focuses on $\Lambda_{\mathbf{u}}(z' | z)$. We can now parameterize $\Lambda_{\mathbf{u}}(z' | z) = \Lambda_{\mathbf{u}}(z' | z; \varphi)$ and optimize its parameterization φ via:

$$\varphi = \operatorname{argmin}_{\varphi} - \mathbb{E}_{(\mathbf{x}', z')} [\log p_{\tau}(z' | \phi(\mathbf{x}'))], \quad (24)$$

$$= \operatorname{argmin}_{\varphi} - \mathbb{E}_{(\mathbf{x}', z')} \left[\log \mathbb{E}_z \left[\Lambda_{\phi(\mathbf{x}')} (z' | z; \varphi) \right] \right] \quad (25)$$

where $(\mathbf{x}', z') \sim (\mathbf{X}^{\tau}, \mathbf{z}^{\tau})$. This will make the target predictor $p_{\tau}(z' | \mathbf{u}) = \mathbb{E}_z[\Lambda_{\mathbf{u}}(z' | z; \varphi)]$ approach the target’s feature-label conditional $D_{\tau}^{\phi}(z' | \mathbf{u})$.

Since $D_{\tau}^{\phi}(z' | \mathbf{u}) = \mathbb{E}_z[\Lambda_{\mathbf{u}}^*(z' | z)]$ (see Definition 5), aligning $p_{\tau}(z' | \mathbf{u})$ with $D_{\tau}^{\phi}(z' | \mathbf{u})$ will decrease the gap between $\Lambda_{\mathbf{u}}(z' | z; \varphi)$ and $\Lambda_{\mathbf{u}}^*(z' | z)$. This will in turn reduce the target fitting term (**TF**) in the target generalization bound in Eq. (10), Theorem 7. See Appendix B for more details.

4 EMPIRICAL ANALYSIS

This section presents a detailed empirical evaluation of our proposed method RECRAFT on two popular benchmarks for cross-modal fine-tuning. These include NASBench-360 (Tu et al., 2022) and PDEBench (Takamoto et al., 2022). NAS-Bench-360 is an extensive benchmark comprising a variety of tasks across 10 distinct data modalities. PDEBench features simulated data from a variety of partial differential equations (PDEs). See Appendix J for additional details.

4.1 Implementation Details

We follow the experiment protocol of ORCA (Shen et al., 2023), using RoBERTa (Liu et al., 2019) for 1D tasks and Swin Transformers (Liu et al., 2021a) for 2D tasks, with CoNLL-2003 and CIFAR-10 as proxy datasets, respectively. Other settings such as learning rates, epochs, and optimizers are taken from ORCA. Full details are provided in Appendix H.

4.2 Results on NAS-Bench-360

The NAS-Bench-360 benchmark (Tu et al., 2022) comprises a diverse set of 10 tasks featuring specialized modalities such as protein sequences, PDE solver, audio, and genetic data, among others. In this set of experiments, we compare 4 types of baselines: (1) hand-designed solution models (Tu et al., 2022); (2) general-purpose models that accepted arbitrary inputs converted to byte arrays (without fine-tuning) such as Perceiver IO (Jaegle et al., 2021); (3) neural architecture search (NAS) methods (no knowledge transfer of existing pre-trained FMs) such as DASH (Shen et al., 2022); and (4) fine-tuning approaches including naive fine-tuning (NFT) and cross-modal fine-tuning methods such as ORCA (Shen et al., 2023), PARE (Cai et al., 2024), MoNA (Ma et al., 2024), and our proposed method RECRAFT.

Table 2 reports the prediction errors achieved by the above baselines across 10 diverse tasks on the NAS-Bench-360 benchmark. It can be observed that RECRAFT achieves the lowest prediction error rates on 8 out of 10 tasks, and second lowest error rate on 1 task. RECRAFT also achieves the best average rank among all baselines. This demonstrates RECRAFT’s robustness in bridging the semantic gap between source and target tasks. We also provide the ablations across several tasks isolating contributions: NFT vs. FA-only vs. RECRAFT, as shown in Table 7 (Appendix E), RECRAFT achieves the best performance across all tasks. Overall, the result underscores the importance of accounting for feature-label distortion during representation alignment (Definition 5) which was overlooked in previous work.

4.3 Results on PDEBench

PDEBench comprises multiple scientific datasets with simulated data from a wide variety of partial differential equations (PDEs) in physics. We compare the fine-tuning performance of RECRAFT with those of naive fine-tuning (NFT) and prior work on cross-modal fine-tuning methods such as ORCA (Shen et al., 2023), PARE (Cai et al., 2024), and MoNA (Ma et al., 2024). Table 1 reports the prediction error achieved by

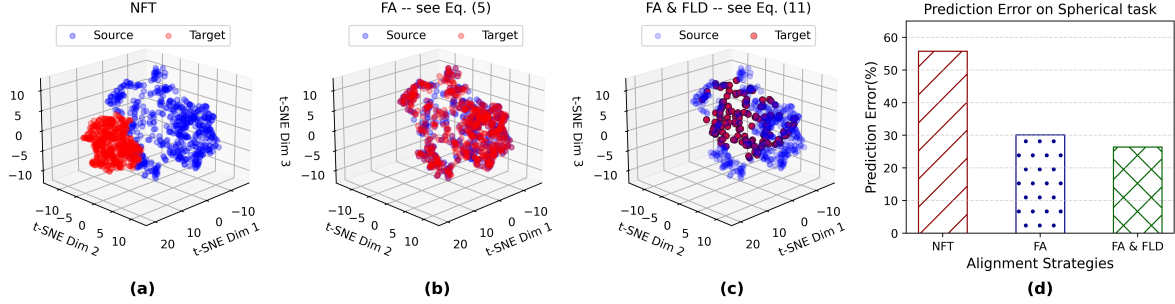


Figure 2: Visualizations of representation alignment (via tSNE) under 3 settings: (a) naive fine-tuning (NFT), which ignores alignment; (b) minimization of feature alignment (**FA**) via Eq. (5); and (c) minimizing a sum of **FA** and feature-label distortion (**FLD**) via Eq. (11). The corresponding predictive errors are shown in (d). NFT exhibits no alignment, while minimizing **FA** leads to exhaustive alignment. Both results in suboptimal performance. In contrast, minimizing **FA** + **FLD** enables selective alignment and achieves the best performance.

Table 1: Prediction errors ($\downarrow$) incurred by the tested baselines across tasks in PDEBench (Takamoto et al., 2022). The results of MoNA (Ma et al., 2024) are quoted from the corresponding paper due to unavailable code. RECRAFT achieves best performance on 7 out of 8 tasks and overall average ranking of 1.25. The columns **#1** and **#2** report the number of times a method achieve best and second best performance, respectively. See Appendix E for additional details on the error bars.

Model	Darcy (2D) nRMSE	Advection (1D) nRMSE	Burgers (1D) nRMSE	Diffusion- Sorption (1D) nRMSE	Shallow Water (2D) nRMSE	Diffusion- Reaction (2D) nRMSE	Diffusion- Reaction (1D) nRMSE	Navier- Stokes (1D) nRMSE	Avg. Rank	#1	#2
NFT	0.085	0.0140	0.0130	3.1E-3	6.1E-3	0.830	9.2E-3	0.863	5.000	0	0
ORCA	0.081	0.0098	0.0120	1.8E-3	6.0E-3	0.820	3.2E-3	0.066	3.250	0	1
PARE	0.081	0.0032	0.0114	1.9E-3	5.9E-3	0.820	2.9E-3	0.068	3.000	1	5
MoNA	0.079	0.0088	0.0114	1.6E-3	5.7E-3	0.818	2.8E-3	0.054	1.875	3	4
RECRAFT	0.079	0.0078	0.0108	1.6E-3	5.4E-3	0.817	2.8E-3	0.050	1.250	7	1

the above baselines across all tasks. It is observed that RECRAFT performs best in 7/8 tasks and second best in the remaining task. It thus achieves the best overall average rank of 1.25. This is consistent with our earlier observation on the NAS-Bench-360 benchmark (Tu et al., 2022), which interestingly demonstrates the effective physic-tuning capability of RECRAFT. It also outperforms existing specialized physic-informed methods such as Fourier neural operators (Li et al., 2021b) in 4/8 tasks as shown in Table 3 in Appendix E.

4.4 Impact of Minimizing Semantic Gap

This section presents additional empirical evidence to support the insight in Theorem 7 that minimizing the **semantic gap** in Eq. (11) (Section 3.1) tightens the upper bound on the generalized target error, thereby improving cross-modal fine-tuning performance. To validate this connection, we track the prediction error and the corresponding semantic gap across optimization iterations used to learn the feature map in stage 1 (see Section 3.1) on three representative tasks (ECG, NinaPro, and DeepSEA) from NAS-Bench-360. As shown in Fig. 3, we observe a strong positive correlation between semantic gap and prediction error, with Pearson correlation coefficients of 0.996 (ECG), 0.965 (NinaPro), and 0.989 (DeepSEA). These results high-

light the practical relevance of the theoretical bound in Theorem 7 and affirm the value of semantic gap minimization as an effective principle for representation learning in cross-modal fine-tuning.

We also study the effect of incorporating feature-label distortion (see Definition 5) in the semantic gap formulation of Eq. (11). This is illustrated in Fig. 2, which compares source-target representation alignment under different alignment strategies. Fig. 2a and Fig. 2b show tSNE feature embeddings under exclusive feature alignment (see Definition 4) and naive fine-tuning (NFT). NFT produces no alignment with the source’s representation space, while exclusive feature alignment leads to an exhaustive alignment. In contrast, Fig. 2c shows that minimizing a combination of feature alignment and feature-label distortion leads to a more selective and effective alignment: target features align only with relevant regions of the source’s space. This leads to substantial performance gains over both NFT and exclusive feature alignment (Table 7, Appendix E).

4.5 Systematic Analysis of Existing Methods

The measurable terms **FA** and **FLD** in the theoretical bound on the target generalization error in Theorem 7 provide a principled diagnostic tool for understanding the strengths and failure modes of prior cross-modal

Table 2: Prediction errors ($\downarrow$) incurred by the tested baselines across 10 diverse tasks in NAS-Bench-360 (Tu et al., 2022). The results of MoNA (Ma et al., 2024) are quoted from the corresponding paper due to unavailable code. RECRAFT achieves best performance on 8 out of 10 tasks which results in the best overall average rank of 1.3 across all tasks. The columns **#1** and **#2** report the number of times a method achieve best and second best performance, respectively. See Appendix E for additional details on the error bars.

Model	Darcy Relative ℓ_2	DeepSEA 1- AUROC	ECG 1-F ₁ score	CIFAR100 0-1 error (%)	Satellite 0-1 error (%)	Spherical 0-1 error(%)	Ninapro 0-1 error (%)	Cosmic 1- AUROC	Psicov MAE _s	FSD50K 1-mAP	Avg. Rank	#1	#2
Hand-designed	8.0E-3	0.30	0.28	19.39	19.80	67.41	8.74	0.13	3.37	0.62	5.6	0	1
NAS-Bench-360	2.6E-2	0.32	0.34	23.39	12.51	48.23	7.35	0.23	2.95	0.63	5.8	0	0
DASH	8.0E-3	0.28	0.32	24.37	12.28	71.38	6.63	0.20	3.30	0.60	4.9	0	2
Perceiver IO	2.4E-2	0.38	0.66	70.04	15.96	82.57	22.4	0.49	8.10	0.73	7.7	0	0
NFT	7.4E-3	0.490	0.44	9.74	13.82	55.76	8.35	0.17	1.92	0.63	5.6	0	0
ORCA	7.5E-3	0.291	0.30	7.80	11.63	29.87	7.74	0.15	1.91	0.56	3.8	0	2
PaRE	7.4E-3	0.286	0.28	6.70	11.21	27.04	7.12	0.12	0.99	0.55	2.2	2	4
MoNA	6.8E-3	0.280	0.27	6.48	11.13	27.13	7.28	0.121	0.99	0.55	1.9	5	2
RECRAFT	7.2E-3	0.278	0.27	7.30	11.11	26.41	6.60	0.11	0.99	0.55	1.3	8	1

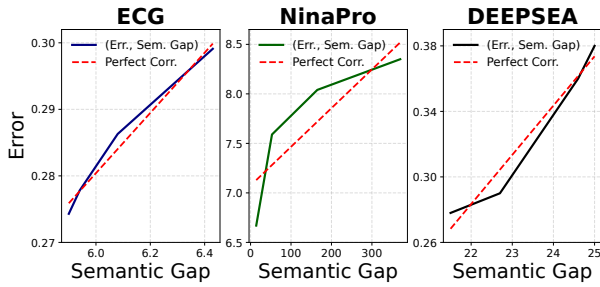


Figure 3: Target error versus semantic gap (Eq. (11)) for RECRAFT on ECG, NinaPro, and DeepSEA. All plots show a strong, consistent correlation across tasks.

adaptation techniques, which remain largely empirical. This section provides additional experiments to measure the **FA** (Eq. 14), **FLD** (Eq. 7), and **TF** (Eq. 9) terms incurred by previous methods (Fig. 4).

Based on the results illustrated in Fig. 4, we can draw the following conclusions. First, NFT (Naive Fine-Tuning) incurs much larger **FA** and **FLD** losses than the other methods and consequently achieves the worst performance (see Tab. 2 and Tab. 1). Second, ORCA effectively reduces **FA** but barely reduces **FLD** by an insignificant amount. This is not surprising given that ORCA performs **FA** and **TF** in separate stages without accounting for their incurred **FLD**. This can make **TF** overcompensate for a suboptimal **FA**, which increases the risk of overfitting. Consequently, its effectiveness remains limited compared to our approach.

Third, PARE has a better balance between **FA** and **FLD** compared to ORCA. However, its significant reduction on **FLD** appears to come at a cost of increasing **FA** (compared to ORCA), which leads to an overall marginal reduction in **FA** + **FLD**. The significant reduction on **FLD** can be attributed to PARE’s design which aims to minimize a linear combination of source and target losses. It uses an intermediate distributional representation that combines the most important information from both source and target modal-

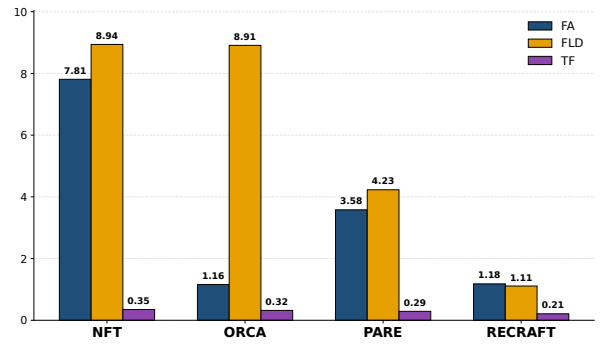


Figure 4: Detailed breakdown of the performance bound in Theorem 7 into **FA**, **FLD**, and **TF** achieved with different cross-modal fine-tuning methods on the Cosmic dataset. The results show that RECRAFT incurs much lower **FLD** than other baselines.

ities. Including most important information from the source modalities in the aligned representation helps constrain the feature alignment (**FA**) within the most important regions in the source’s representation landscape, which in turn limits how far the aligned features can deviate from the source’s feature-label structure. This can help reduce **FLD**. Nonetheless, this scheme was not optimized to prevent negative transfer, where some source-specific important information is irrelevant to the target’s task but might still be included in the intermediate representation (see Fig. 2). As such, the feature alignment must account for such extra, irrelevant information and hence, increases the **FA** loss. This explains why PARE incurs larger **FA** than ORCA despite having smaller **FLD**.

In contrast, our method RECRAFT implements a principled minimization procedure for both **FA** + **FLD** (Section 3.1) and **TF** (Section 3.2). This leads to a significant reduction in both **FA** and **FLD** while also reducing **TF**, thus consistently achieving better performance than all the above baselines (Section 4). Moreover, our theoretical decomposition provides a new analytical lens that will inspire important

and synergistic research directions (Appendix K).

5 RELATED WORKS

Due to limited space, we provide a short review of the most relevant work on in-modal and cross-modal fine-tuning. A broader review is deferred to Appendix C.

5.1 In-Modal Fine-Tuning of FMs

The recent advent of large pre-trained or foundation models (FMs) has enabled flexible knowledge transfer to a wide range of downstream tasks under a unified, task-agnostic paradigm. This is commonly known as fine-tuning, which is model-agnostic and does not require shared input or output spaces. Such transfer is feasible because these FMs are trained on broad, diverse corpora (e.g., text, image, or audio) that often overlap semantically or structurally with the content of downstream datasets for in-modal tasks. As a result, fine-tuning has been successfully applied across domains such as vision (Kirillov et al., 2023; Liu et al., 2021b; Bui et al., 2024; Weng et al., 2024; Phung et al., 2026), video understanding (Bertasius et al., 2021), language (Zheng et al., 2021; Yang et al., 2022; Ma et al., 2023), and speech (Radford et al., 2022; Li et al., 2021a). There are also multi-modal FMs (Radford et al., 2021; Alayrac et al., 2022; Kim et al., 2021) which were pre-trained to learn the embeddings of multiple modalities together but existing approaches to facilitate knowledge transfer from these models are still restricted to within the pre-trained modalities.

5.2 Cross-Modal Fine-Tuning of FMs

Early attempts in cross-modal fine-tuning have focused on transferring language models to other modalities such as vision (Kiela et al., 2019; Tan and Bansal, 2019; Gu et al., 2022), DNA/protein sequences (Nguyen et al., 2023; Lin et al., 2023; Jumper et al., 2021). These provide initial evidence of cross-modal transferability but their designs are hand-tailored to specific target tasks and modalities rather than for general purpose (Shen et al., 2023). Recently, a few general-purpose cross-modal fine-tuning framework have emerged with remarkable successes in finetuning existing vision (Liu et al., 2021b) or language (Liu et al., 2019) FMs to solve a broad, diverse set of unseen tasks and data modalities (Shen et al., 2023; Cai et al., 2024; Ma et al., 2024).

In particular, ORCA (Shen et al., 2023) aligns source and target representation distributions using optimal transport before fine-tuning the entire network. PARE (Cai et al., 2024) alternatively introduces a gating mechanism to integrate source and target features during fine-tuning. MoNA (Ma et al., 2024)

addresses modality representation alignment via a bi-level optimization framework: the inner loop learns the optimal target predictor for a given embedder, while the outer loop updates the embedder to align the source feature-label semantics with the combined representation and prediction under this predictor. This approach heuristically reduces the misalignment between source’s and target’s feature-label structures while fitting to the target data. However, as noted in Section 1, these approaches largely adopt heuristic methods combining distributional feature alignment and target fitting to facilitate cross-modal knowledge transfer. Overall, these approaches lack a principled means to assess the impact of their heuristic strategies on the generalized target performance.

6 LIMITATIONS AND FUTURE WORKS

In future work, our proposed method will be further generalized to address two existing limitations. First, while the use of Euclidean cost metric to characterize the geometric structure under which the Wasserstein distance is computed does not invalidate the bound in Theorem 7, it is unclear whether a Wasserstein distance parameterized with Euclidean cost metric would optimally tighten the upper-bound. We will investigate this aspect via developing a metric learning component which parameterizes and learns this cost metric as additional parameters of the upper-bound. Second, while upper-bounding **FLD** with the conditional entropic terms in Eq. 18 and approximating them with a well-established pseudo-label approach (Nguyen et al., 2020) results in a valid empirical upper-bound, it is unclear whether the bound gap can be made vanishingly small as we increase the no. of sampled pseudo labels (see Eq. 19). This remains an open question to be addressed in the future generalization of our work.

7 CONCLUSION

This paper presents a new theoretical perspective for cross-modal fine-tuning which highlights and addresses a limitation that hinders generalized knowledge transfer in existing work. This leads to a novel development of a theoretical framework that provably establishes an upper-bound on the generalized target performance. The bound reveals a precise computational representation for the interaction between feature alignment and target fitting. This inspires a practical algorithm that optimize this interaction in a principled manner, paving the way for more effective algorithm designs. The developed algorithm performs significantly better than the recent SOTA methods on two extensive cross-modal fine-tuning benchmark, thus conclusively validating our theoretical insight.

Acknowledgement

This work utilized GPU compute resource at SDSC and ACES through allocation CIS230391 from the Advanced Cyberinfrastructure Coordination Ecosystem: Services and Support (ACCESS) program Boerner et al. (2023), which is supported by U.S. National Science Foundation grants #2138259, #2138286, #2138307, #2137603, and #2138296.

References

- Jean-Baptiste Alayrac, Jeff Donahue, Pauline Luc, Antoine Miech, Iain Barr, Yana Hasson, Karel Lenc, Arthur Mensch, Katie Millican, Malcolm Reynolds, Roman Ring, Eliza Rutherford, Serkan Cabi, Tengda Han, Zhitao Gong, Sina Samangooei, Marianne Monteiro, Jacob Menick, Sebastian Borgeaud, Andy Brock, Aida Nematzadeh, Sahand Sharifzadeh, Mikolaj Binkowski, Ricardo Barreira, Oriol Vinyals, Andrew Zisserman, and Karen Simonyan. Flamingo: a visual language model for few-shot learning. *ArXiv*, abs/2204.14198, 2022. URL <https://api.semanticscholar.org/CorpusID:248476411>.
- Gedas Bertasius, Heng Wang, and Lorenzo Torresani. Is space-time attention all you need for video understanding?, 2021.
- Timothy J. Boerner, Stephen Deems, Thomas R. Furlani, Shelley L. Knuth, and John Towns. Access: Advancing innovation: Nsf’s advanced cyberinfrastructure coordination ecosystem: Services & support. In *Practice and Experience in Advanced Research Computing 2023: Computing for the Common Good*, PEARC ’23, page 173–176, New York, NY, USA, 2023. Association for Computing Machinery. ISBN 9781450399852. doi: 10.1145/3569951.3597559. URL <https://doi.org/10.1145/3569951.3597559>.
- Long Minh Bui, Tho Tran Huu, Duy Dinh, Tan Minh Nguyen, and Trong Nghia Hoang. Revisiting kernel attention with correlated gaussian process representation. In *The 40th Conference on Uncertainty in Artificial Intelligence*, 2024. URL <https://openreview.net/forum?id=xLIK0vu3MW>.
- Lincan Cai, Shuang Li, Wenxuan Ma, Jingxuan Kang, Binhui Xie, Zixun Sun, and Chengwei Zhu. Enhancing cross-modal fine-tuning with gradually intermediate modality generation, 2024. URL <https://arxiv.org/abs/2406.09003>.
- Thomas M. Cover and Joy A. Thomas. *Elements of Information Theory*. Wiley-Interscience, 2nd edition, July 2006. Theorem 2.6.5 states “Conditioning reduces entropy ($H(X-Y) \leq H(X)$), sometimes summarized as ‘information never hurts.’”.
- Junnan Gu, Xi Chen, Yang Liu, and Ming Li. Vision-language pre-training with triple contrastive learning, 2022.
- Minh Hoang and Trong Nghia Hoang. Few-shot learning via repurposing ensemble of black-box models. In *Proc. AAAI*, 2024.
- Trong Nghia Hoang, Chi Thanh Lam, Kian Hsiang Low, and Patrick Jaillet. Learning Task-Agnostic Embedding of Multiple Black-box Experts for Multi-Task Model Fusion. In *Proc. ICML*, 2020.
- Andrew Jaegle, Sebastian Borgeaud, Jean-Baptiste Alayrac, Carl Doersch, Catalin Ionescu, David Ding, Skanda Koppula, Daniel Zoran, Andrew Brock, Evan Shelhamer, et al. Perceiver io: A general architecture for structured inputs & outputs. *arXiv preprint arXiv:2107.14795*, 2021.
- John Jumper, Richard Evans, Alexander Pritzel, Tim Green, Michael Figurnov, Olaf Ronneberger, Kathryn Tunyasuvunakool, Russ Bates, Augustin Žídek, Anna Potapenko, Alex Bridgland, Clemens Meyer, Simon A A Kohl, Andrew J Ballard, Andrew Cowie, Bernardino Romera-Paredes, Stanislav Nikolov, Rishub Jain, Jonas Adler, Trevor Back, Stig Petersen, David Reiman, Ellen Clancy, Michal Zielinski, Martin Steinegger, Michalina Pacholska, Tamas Berghammer, Sebastian Bodenstein, David Silver, Oriol Vinyals, Andrew W Senior, Koray Kavukcuoglu, Pushmeet Kohli, and Demis Hassabis. Highly accurate protein structure prediction with alphafold, 2021.
- Douwe Kiela, Suvrat Bhooshan, Hamed Firooz, and Davide Testuggine. Supervised multimodal bitransformers for classifying images and text. *ArXiv*, abs/1909.02950, 2019. URL <https://api.semanticscholar.org/CorpusID:202539204>.
- Wonjae Kim, Bokyung Son, and Ildoo Kim. Vilt: Vision-and-language transformer without convolution or region supervision. In *International Conference on Machine Learning*, 2021. URL <https://api.semanticscholar.org/CorpusID:231839613>.
- Diederik P. Kingma and Jimmy Ba. Adam: A method for stochastic optimization, 2017. URL <https://arxiv.org/abs/1412.6980>.
- Alexander Kirillov, Eric Mintun, Nikhila Ravi, Hanzi Mao, Chloe Rolland, Laura Gustafson, Tete Xiao, Spencer Whitehead, Alexander C Berg, Wan-Yen Lo, Piotr Dollár, and Ross Girshick. Segment anything, 2023.
- Chi Thanh Lam, Trong Nghia Hoang, Kian Hsiang Low, and Patrick Jaillet. Model Fusion for Personalized Learning. In *Proc. ICML*, 2021.
- Shuang Li, Binhui Xie, Jiashu Wu, Ying Zhao, Chi Harold Liu, and Zhengming Ding. Simul-

- taneous semantic alignment network for heterogeneous domain adaptation. *Proceedings of the 28th ACM International Conference on Multimedia*, 2020. URL <https://api.semanticscholar.org/CorpusID:220961739>.
- Xian Li, Changan Wang, Yun Tang, Chau Tran, Yuqing Tang, Juan Pino, Alexei Baevski, Alexis Conneau, and Michael Auli. Multilingual speech translation from efficient finetuning of pretrained models. In Chengqing Zong, Fei Xia, Wenjie Li, and Roberto Navigli, editors, *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 827–838, Online, August 2021a. Association for Computational Linguistics. doi: 10.18653/v1/2021.acl-long.68. URL <https://aclanthology.org/2021.acl-long.68/>.
- Zongyi Li, Nikola Borislavov Kovachki, Kamyar Azizzadenesheli, Burigede liu, Kaushik Bhattacharya, Andrew Stuart, and Anima Anandkumar. Fourier neural operator for parametric partial differential equations. In *International Conference on Learning Representations*, 2021b. URL <https://openreview.net/forum?id=c8P9NQVtmn0>.
- Yuxiang Lin, Ling Luo, Ying Chen, Xushi Zhang, Zihui Wang, Wenxian Yang, Mengsha Tong, and Rongshan Yu. St-align: A multimodal foundation model for image-gene alignment in spatial transcriptomics. *ArXiv*, abs/2411.16793, 2024. URL <https://api.semanticscholar.org/CorpusID:274280682>.
- Zeming Lin, Halil Akin, Roshan Rao, Brian Hie, Zhongkai Zhu, Wenting Lu, Nikita Smetanin, Robert Verkuil, Ori Kabeli, Yaniv Shmueli, Allan dos Santos Costa, Maryam Fazel-Zarandi, Tom Sercu, Salvatore Candido, and Alexander Rives. Evolutionary-scale prediction of atomic-level protein structure with a language model. *Science*, 379(6637):1123–1130, 2023. doi: 10.1126/science.ade2574. URL <https://www.science.org/doi/abs/10.1126/science.ade2574>. Earlier versions as preprint: bioRxiv 2022.07.20.500902.
- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. Roberta: A robustly optimized bert pretraining approach. *ArXiv*, abs/1907.11692, 2019.
- Ze Liu, Yutong Lin, Yue Cao, Han Hu, Yixuan Wei, Zheng Zhang, Stephen Lin, and Baining Guo. Swin transformer: Hierarchical vision transformer using shifted windows. *2021 IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 9992–10002, 2021a.
- Ze Liu, Yutong Lin, Yue Cao, Han Hu, Yixuan Wei, Zheng Zhang, Stephen Lin, and Baining Guo. Swin transformer: Hierarchical vision transformer using shifted windows, 2021b.
- Ilya Loshchilov, Frank Hutter, et al. Fixing weight decay regularization in adam. *arXiv preprint arXiv:1711.05101*, 5(5):5, 2017.
- Wenxuan Ma, Shuang Li, Lincan Cai, and Jingxuan Kang. Language semantic graph guided data-efficient learning. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023. URL <https://openreview.net/forum?id=tUyW68cRqr>.
- Wenxuan Ma, Shuang Li, Lincan Cai, and Jingxuan Kang. Learning modality knowledge alignment for cross-modality transfer, 2024. URL <https://arxiv.org/abs/2406.18864>.
- Cuong V. Nguyen, Tal Hassner, Matthias Seeger, and Cedric Archambeau. Leep: A new measure to evaluate transferability of learned representations. In *Proceedings of the 37th International Conference on Machine Learning*, pages 7294–7305. PMLR, 2020.
- Eric Nguyen, Michael Poli, Marjan Faizi, Armin Thomas, Callum Birch-Sykes, Michael Wornow, Aman Patel, Clayton Rabideau, Stefano Massaroli, Yoshua Bengio, Stefano Ermon, Stephen A. Baccus, and Chris Ré. Hyenadna: Long-range genomic sequence modeling at single nucleotide resolution, 2023. URL <https://arxiv.org/abs/2306.15794>.
- Sinno Jialin Pan and Qiang Yang. A survey on transfer learning. *IEEE Transactions on Knowledge and Data Engineering*, 22(10):1345–1359, 2010. doi: 10.1109/TKDE.2009.191.
- Gabriel Peyré and Marco Cuturi. Computational optimal transport, 2020. URL <https://arxiv.org/abs/1803.00567>.
- Thu Hang Phung, Duong M. Nguyen, Thanh Trung Huynh, Quoc Viet Hung Nguyen, Trong Nghia Hoang, and Phi Le Nguyen. Federated prompt-tuning with heterogeneous and incomplete multimodal client data. *ArXiv*, abs/2602.07081, 2026. URL <https://api.semanticscholar.org/CorpusID:284585890>.
- Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, Gretchen Krueger, and Ilya Sutskever. Learning transferable visual models from natural language supervision. In *International Conference on Machine Learning*, 2021. URL <https://api.semanticscholar.org/CorpusID:231591445>.
- Alec Radford, Jong Wook Kim, Tao Xu, Greg Brockman, Christine McLeavey, and Ilya Sutskever. Ro-

- bust speech recognition via large-scale weak supervision, 2022. URL <https://arxiv.org/abs/2212.04356>.
- M. Raissi, P. Perdikaris, and G.E. Karniadakis. Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations. *Journal of Computational Physics*, 378:686–707, 2019. ISSN 0021-9991. doi: <https://doi.org/10.1016/j.jcp.2018.10.045>. URL <https://www.sciencedirect.com/science/article/pii/S0021999118307125>.
- Olaf Ronneberger, Philipp Fischer, and Thomas Brox. U-net: Convolutional networks for biomedical image segmentation. In Nassir Navab, Joachim Hornegger, William M. Wells, and Alejandro F. Frangi, editors, *Medical Image Computing and Computer-Assisted Intervention – MICCAI 2015*, pages 234–241, Cham, 2015. Springer International Publishing.
- Jian Shen, Yanru Qu, Weinan Zhang, and Yong Yu. Wasserstein distance guided representation learning for domain adaptation, 2018. URL <https://arxiv.org/abs/1707.01217>.
- Junhong Shen, Mikhail Khodak, and Ameet Talwalkar. Efficient architecture search for diverse tasks, 2022. URL <https://arxiv.org/abs/2204.07554>.
- Junhong Shen, Liam Li, Lucio M. Dery, Corey Staten, Mikhail Khodak, Graham Neubig, and Ameet Talwalkar. Cross-modal fine-tuning: Align then refine. In *Proceedings of the International Conference on Machine Learning (ICML)*. ICML, 2023. URL <https://arxiv.org/abs/2302.05738>.
- Makoto Takamoto, Timothy Praditia, Raphael Leteritz, Dan MacKinlay, Francesco Alesiani, Dirk Pflüger, and Mathias Niepert. Pdebench: An extensive benchmark for scientific machine learning. In *Advances in Neural Information Processing Systems (NeurIPS) Datasets and Benchmarks Track*, 2022.
- Hao Tan and Mohit Bansal. Lxmert: Learning cross-modality encoder representations from transformers, 2019.
- Renbo Tu, Nicholas Roberts, Mikhail Khodak, Junhong Shen, Frederic Sala, and Ameet Talwalkar. NAS-bench-360: Benchmarking neural architecture search on diverse tasks. In *Thirty-sixth Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2022. URL <https://openreview.net/forum?id=xUXTbq6gWsB>.
- C. Villani. *Optimal Transport: Old and New*. Grundlehren der mathematischen Wissenschaften. Springer Berlin Heidelberg, 2008. ISBN 9783540710509. URL https://books.google.com.vn/books?id=hV8o5R7_5tkC.
- Mei Wang and Weihong Deng. Deep visual domain adaptation: A survey. *Neurocomputing*, 312:135–153, 2018. ISSN 0925-2312. doi: <https://doi.org/10.1016/j.neucom.2018.05.083>. URL <https://www.sciencedirect.com/science/article/pii/S0925231218306684>.
- Pei-Yau Weng, Minh Hoang, Lam M. Nguyen, My T. Thai, Tsui-Wei Weng, and Trong Nghia Hoang. Probabilistic federated prompt-tuning with non-IID and imbalanced data. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024. URL <https://openreview.net/forum?id=nw6ANsC66G>.
- Huiyun Yang, Huadong Chen, Hao Zhou, and Lei Li. Enhancing cross-lingual transfer by manifold mixup. In *International Conference on Learning Representations*, 2022. URL <https://openreview.net/forum?id=0jPmfr9GkVv>.
- Yuan Yao, Yu Zhang, Xutao Li, and Yunming Ye. Heterogeneous domain adaptation via soft transfer network. *Proceedings of the 27th ACM International Conference on Multimedia*, 2019. URL <https://api.semanticscholar.org/CorpusID:201651510>.
- Wenzhe Yin, Shujian Yu, Yicong Lin, Jie Liu, Jan-Jakob Sonke, and Efstratios Gavves. Domain adaptation with cauchy-schwarz divergence, 2024. URL <https://arxiv.org/abs/2405.19978>.
- Bo Zheng, Li Dong, Shaohan Huang, Wenhui Wang, Zewen Chi, Saksham Singhal, Wanxiang Che, Ting Liu, Xia Song, and Furu Wei. Consistency Regularization for Cross-Lingual Fine-Tuning. In *Proceedings of ACL 2021*, 2021.

Checklist

1. For all models and algorithms presented, check if you include:
 - (a) A clear description of the mathematical setting, assumptions, algorithm, and/or model. [Yes]
 - (b) An analysis of the properties and complexity (time, space, sample size) of any algorithm. [Yes]
 - (c) (Optional) Anonymized source code, with specification of all dependencies, including external libraries. [Yes]
2. For any theoretical claim, check if you include:
 - (a) Statements of the full set of assumptions of all theoretical results. [Yes]
 - (b) Complete proofs of all theoretical results. [Yes]
 - (c) Clear explanations of any assumptions. [Yes]
3. For all figures and tables that present empirical results, check if you include:
 - (a) The code, data, and instructions needed to reproduce the main experimental results (either in the supplemental material or as a URL). [Yes, see Appendix H]
 - (b) All the training details (e.g., data splits, hyperparameters, how they were chosen). [Yes]
 - (c) A clear definition of the specific measure or statistics and error bars (e.g., with respect to the random seed after running experiments multiple times). [Yes]
 - (d) A description of the computing infrastructure used. (e.g., type of GPUs, internal cluster, or cloud provider). [Yes, see Appendix H]
4. If you are using existing assets (e.g., code, data, models) or curating/releasing new assets, check if you include:
 - (a) Citations of the creator If your work uses existing assets. [Yes]
 - (b) The license information of the assets, if applicable. [Not Applicable]
 - (c) New assets either in the supplemental material or as a URL, if applicable. [Yes]
 - (d) Information about consent from data providers/curators. [Not Applicable, Public datasets]
 - (e) Discussion of sensible content if applicable, e.g., personally identifiable information or offensive content. [Not Applicable]
5. If you used crowdsourcing or conducted research with human subjects, check if you include:
 - (a) The full text of instructions given to participants and screenshots. [Not Applicable]
 - (b) Descriptions of potential participant risks, with links to Institutional Review Board (IRB) approvals if applicable. [Not Applicable]
 - (c) The estimated hourly wage paid to participants and the total amount spent on participant compensation. [Not Applicable]

Supplementary Materials

A Proof of Theorem 7

This section provides the detailed proof of our main result stated in Theorem 7, which bounds the gap between the target and source generalized errors. The bound is characterized in terms of (i) the distributional feature alignment (**FA**) between the source and target in Eq. (5), (ii) the feature-label distortion (**FLD**) under their respective feature maps in Eq. (7), and (iii) the target model's fit (**TF**) to the fine-tuning data in Eq. (9). For clarity, we restate the result below.

$$\text{err}_\tau(\phi) \leq \text{err}_s(\theta) + \mathbf{FA}(\phi, \theta) + \mathbb{E}_{D_\tau^\phi(\mathbf{u})}[\mathbf{FLD}(\mathbf{u}) + \mathbf{TF}(\mathbf{u})], \quad (26)$$

Here, $D_\tau^\phi(\mathbf{u})$ denote the target's marginal feature distribution under (target) feature map ϕ and $D_s^\theta(z | \mathbf{u})$ denote the feature-label conditional under (source) feature map θ . Our proof goes below.

First, following Definition 1, the generalized target error is

$$\text{err}_\tau(\phi) \triangleq -\mathbb{E}_{(\mathbf{x}', z') \sim D_\tau}[\log p_\tau(z' | \phi(\mathbf{x}'))] \quad (27)$$

$$= -\mathbb{E}_{(\mathbf{u}, z') \sim D_\tau^\phi}[\log p_\tau(z' | \mathbf{u})] = -\mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_\tau^\phi(z' | \mathbf{u})}[\log p_\tau(z' | \mathbf{u})]. \quad (28)$$

Likewise, the generalized source error is

$$\text{err}_s(\theta) = -\mathbb{E}_{D_s^\theta(\mathbf{u})} \mathbb{E}_{D_s^\theta(z | \mathbf{u})}[\log p_s(z | \mathbf{u})]. \quad (29)$$

Combining Eqs. (28) and (29), we can rewrite

$$\text{err}_\tau(\phi) - \text{err}_s(\theta) = \mathbf{A} + \mathbf{B} \quad \text{where} \quad (30)$$

$$\mathbf{A} \triangleq \text{err}_\tau(\phi) + \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z | \mathbf{u})} \log D_s^\theta(z | \mathbf{u}), \quad (31)$$

$$\mathbf{B} \triangleq -\text{err}_s(\theta) - \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z | \mathbf{u})} \log D_s^\theta(z | \mathbf{u}). \quad (32)$$

We will bound $\mathbf{A}$ and $\mathbf{B}$ next.

1. Bounding A. Plugging Eq. (28) into Eq. (31), we have

$$\mathbf{A} = \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z | \mathbf{u})}[\log D_s^\theta(z | \mathbf{u})] - \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_\tau^\phi(z' | \mathbf{u})}[\log p_\tau(z' | \mathbf{u})]. \quad (33)$$

Following Definition 5, let $\Lambda_{\mathbf{u}}^*(z' | z) \in C_{\mathbf{u}}^*$ denote a valid transport map from $D_s^\theta(z | \mathbf{u})$ to $D_\tau^\phi(z' | \mathbf{u})$. That is,

$$D_\tau^\phi(z' | \mathbf{u}) = \mathbb{E}_{D_s^\theta(z | \mathbf{u})}[\Lambda_{\mathbf{u}}^*(z' | z)]. \quad (34)$$

Plugging Eq. (34) into Eq. (33), we can rewrite

$$\mathbf{A} = \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z | \mathbf{u})}[\log D_s^\theta(z | \mathbf{u})] - \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z | \mathbf{u})} \left[\sum_{z'} \Lambda_{\mathbf{u}}^*(z' | z) \log p_\tau(z' | \mathbf{u}) \right]. \quad (35)$$

Furthermore, we also have

$$\begin{aligned} -\log p_\tau(z' | \mathbf{u}) &= -\log \left(\sum_a \Lambda_{\mathbf{u}}(z' | a) D_s^\theta(a | \mathbf{u}) \right) \leq -\log \left(\Lambda_{\mathbf{u}}(z' | z) D_s^\theta(z | \mathbf{u}) \right) \\ &= -\log \Lambda_{\mathbf{u}}(z' | z) - \log D_s^\theta(z | \mathbf{u}). \end{aligned} \quad (36)$$

for any z and valid transport map $\Lambda_{\mathbf{u}}(z' | \cdot) \in C_{\mathbf{u}}$ from $D_s^\theta(z | \mathbf{u})$ to $p_\tau(z' | \mathbf{u})$ as defined in Definition 6. Plugging Eq. (36) into Eq. (35), we obtain the following upper-bound,

$$\begin{aligned}
 \mathbf{A} &\leq \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[\log D_s^\theta(z | \mathbf{u}) \right] \\
 &\quad - \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[\sum_{z'} \Lambda_{\mathbf{u}}^*(z' | z) \log \Lambda_{\mathbf{u}}(z' | z) \right] - \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[\sum_{z'} \Lambda_{\mathbf{u}}^*(z' | z) \log D_s^\theta(z | \mathbf{u}) \right] \quad (37) \\
 &= \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[\log D_s^\theta(z | \mathbf{u}) \right] \\
 &\quad - \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[\sum_{z'} \Lambda_{\mathbf{u}}^*(z' | z) \log \Lambda_{\mathbf{u}}(z' | z) \right] - \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[\log D_s^\theta(z | \mathbf{u}) \right] \\
 &= \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[- \sum_{z'} \Lambda_{\mathbf{u}}^*(z' | z) \log \Lambda_{\mathbf{u}}(z' | z) \right] = \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \mathbb{H} \left[\Lambda_{\mathbf{u}}^*(\cdot | z), \Lambda_{\mathbf{u}}(\cdot | z) \right] \\
 &= \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \mathbb{H} \left[\Lambda_{\mathbf{u}}^*(\cdot | z) \right] + \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \mathbb{KL} \left[\Lambda_{\mathbf{u}}^*(\cdot | z) \parallel \Lambda_{\mathbf{u}}(\cdot | z) \right]. \quad (38)
 \end{aligned}$$

As Eq. (38) holds for all choices of $\Lambda_{\mathbf{u}}^*(z' | z) \in C_{\mathbf{u}}^*$ and $\Lambda_{\mathbf{u}}(z' | z) \in C_{\mathbf{u}}$, we can tighten its right-hand side (RHS) by choosing for each $(\mathbf{u})$, $\Lambda_{\mathbf{u}}^*(\cdot | \cdot) \in C_{\mathbf{u}}^*$ that minimizes $\mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[\mathbb{H}[\Lambda_{\mathbf{u}}^*(\cdot | z)] \right]$ and taking minimum over the remaining choice of $\Lambda_{\mathbf{u}}(\cdot | \cdot) \in C_{\mathbf{u}}$. This results in a tighten bound below:

$$\mathbf{A} \leq \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \left[\min_{\Lambda_{\mathbf{u}}^* \in C_{\mathbf{u}}^*} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[\mathbb{H}[\Lambda_{\mathbf{u}}^*(\cdot | z)] \right] \right] + \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \left[\min_{\Lambda_{\mathbf{u}} \in C_{\mathbf{u}}} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[\mathbb{KL}(\Lambda_{\mathbf{u}}^+(\cdot | z) \parallel \Lambda_{\mathbf{u}}(\cdot | z)) \right] \right]. \quad (39)$$

where $\Lambda_{\mathbf{u}}^+(\cdot | \cdot) = \operatorname{argmin}_{\Lambda_{\mathbf{u}}^* \in C_{\mathbf{u}}^*} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[\mathbb{H}[\Lambda_{\mathbf{u}}^*(\cdot | z)] \right]$. Following the definition of **FLD** and **TF** in Eqs. (7)-(9),

$$\mathbf{A} \leq \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \left[\mathbf{FLD}(\mathbf{u}) + \mathbf{TF}(\mathbf{u}) \right], \quad (40)$$

2. Bounding B. We will now show that **B** in Eq. (32) is upper-bounded by **FA**(ϕ, θ) in Eq. (5) to complete the proof. To see this, note that in practice, $p_s(z | \mathbf{u})$ often approximates $D_s^\theta(z | \mathbf{u})$ faithfully. Exploiting this practical property, we can rewrite **B** in Eq. (32) as

$$\mathbf{B} = \mathbb{E}_{D_\tau^\phi(\mathbf{u})} \left[- \mathbb{E}_{D_s^\theta(z|\mathbf{u})} [\log p_s(z | \mathbf{u})] \right] - \mathbb{E}_{D_s^\theta(\mathbf{u})} \left[- \mathbb{E}_{D_s^\theta(z|\mathbf{u})} [\log p_s(z | \mathbf{u})] \right] \quad (41)$$

$$= \mathbb{E}_{D_\tau^\phi(\mathbf{u})} [\ell_s(\mathbf{u})] - \mathbb{E}_{D_s^\theta(\mathbf{u})} [\ell_s(\mathbf{u})]. \quad (42)$$

where $\ell_s(\mathbf{u}) \triangleq \mathbb{E}_{D_s^\theta(z|\mathbf{u})} [-\log p_s(z | \mathbf{u})]$ as previously defined in Definition 4. For any cost metric $\delta \in \Delta$ such that $|\ell_s(\mathbf{u}_1) - \ell_s(\mathbf{u}_2)| \leq \tau_\delta \cdot \delta(\mathbf{u}_1, \mathbf{u}_2)$, the Kantorovich-Rubinstein duality ascertains that

$$\mathbf{B} \leq \tau_\delta \cdot W_\delta(D_\tau^\phi(\mathbf{u}), D_s^\theta(\mathbf{u})). \quad (43)$$

As this is true for any $\delta \in \Delta$, we can again tighten the right-hand side (RHS) of Eq. (43) via taking minimum over the choice of δ . That is,

$$\mathbf{B} \leq \min_{\delta \in \Delta} \left\{ \tau_\delta \cdot W_\delta(D_\tau^\phi(\mathbf{u}), D_s^\theta(\mathbf{u})) \right\} = \mathbf{FA}(\phi, \theta). \quad (44)$$

Plugging Eqs. (40) and (44) in Eq. (30) completes our proof.

B Tightness of the Generalization Bound in Theorem 7

This section evaluates the empirical tightness of the bound in Theorem 7 which is quoted below for convenience:

$$\operatorname{err}_\tau(\phi) \leq \operatorname{err}_s(\theta) + \mathbf{FA}(\phi, \theta) + \mathbb{E}_{D_\tau^\phi(\mathbf{u})} [\mathbf{FLD}(\mathbf{u}) + \mathbf{TF}(\mathbf{u})]. \quad (45)$$

Table 3: Comparison of RECRAFT and Baselines on Multiple PDE Tasks ($\downarrow$ indicates lower is better). U-Net results for Navier-Stokes and Darcy-Flow are unavailable (due to memory constraints) in the benchmark paper. RECRAFT achieves an average rank of **1.5** and attains the best performance on 4 out of 8 tasks. Columns **#1** and **#2** indicate the number of times each method achieves the best and second-best performance, respectively.

Model	Darcy (2D) nRMSE	Advection (1D) nRMSE	Burgers (1D) nRMSE	Diffusion- Sorption (1D) nRMSE	Shallow Water (2D) nRMSE	Diffusion- Reaction (2D) nRMSE	Diffusion- Reaction (1D) nRMSE	Navier- Stokes (1D) nRMSE	Avr. Rank	#1	#2
PINN	0.18	0.67	0.36	0.15	0.085	0.84	0.84	0.720	3.125	0	1
FNO	0.22	0.011	0.0031	1.8E-3	4.4E-3	0.12	1.4E-3	0.068	1.625	4	3
U-Net	-	1.1	0.99	0.22	0.017	1.6	0.08	-	-	0	0
RECRAFT	0.079	0.0078	0.0108	1.6E-3	5.4E-3	0.817	2.8E-3	0.050	1.500	4	4

Table 4: Prediction errors ($\downarrow$) incurred by the tested methods (with standard deviation) across 10 diverse tasks on NAS-Bench-360. The reported results of MoNA (Ma et al., 2024) are quoted from the corresponding paper since its source code is not released. Our method RECRAFT achieves best performance in 8 out 10 tasks.

Model	Darcy Relative ℓ_2	DeepSEA 1- AUROC	ECG 1-F ₁ score	CIFAR100 0-1 error (%)	Satellite 0-1 error (%)	Spherical 0-1 error (%)	Ninapro 0-1 error (%)	Cosmic 1- AUROC	Psicov MAEs	FSD50K 1-mAP
Hand-designed	8.0E-3 $\pm$ 1E-3	0.300 $\pm$ 2E-2	0.28 $\pm$ 0.01	19.39 $\pm$ 0.2	19.80 $\pm$ 0.01	67.41 $\pm$ 0.8	8.74 $\pm$ 0.9	0.13 $\pm$ 1E-2	3.37 $\pm$ 0.15	0.62 $\pm$ 4E-3
NAS-Bench-360	2.6E-2 $\pm$ 1E-3	0.320 $\pm$ 1E-2	0.34 $\pm$ 0.01	23.39 $\pm$ 0.03	12.51 $\pm$ 0.25	48.23 $\pm$ 2.5	7.35 $\pm$ 0.8	0.23 $\pm$ 4E-3	2.95 $\pm$ 0.14	0.60 $\pm$ 0.03
DASH	8.0E-3 $\pm$ 2E-3	0.280 $\pm$ 1E-2	0.32 $\pm$ 6E-3	24.37 $\pm$ 0.83	12.28 $\pm$ 0.50	71.38 $\pm$ 0.7	6.63 $\pm$ 0.4	0.20 $\pm$ 5E-3	3.30 $\pm$ 0.17	0.60 $\pm$ 0.02
Perceiver IO	2.4E-2 $\pm$ 1E-2	0.380 $\pm$ 4E-3	0.66 $\pm$ 0.01	70.04 $\pm$ 0.4	15.96 $\pm$ 0.01	82.57 $\pm$ 0.2	22.4 $\pm$ 1.5	0.49 $\pm$ 1E-2	8.10 $\pm$ 0.05	0.73 $\pm$ 0.02
NFT	7.4E-3 $\pm$ 1E-4	0.490 $\pm$ 2E-2	0.44 $\pm$ 0.03	9.74 $\pm$ 1.21	13.82 $\pm$ 0.24	55.76 $\pm$ 2.3	8.35 $\pm$ 0.4	0.17 $\pm$ 2E-2	1.92 $\pm$ 0.06	0.63 $\pm$ 0.01
ORCA	7.5E-3 $\pm$ 8E-5	0.291 $\pm$ 2E-3	0.30 $\pm$ 6E-3	7.80 $\pm$ 0.41	11.63 $\pm$ 0.20	29.87 $\pm$ 0.8	7.74 $\pm$ 0.4	0.15 $\pm$ 5E-3	1.91 $\pm$ 0.04	0.56 $\pm$ 0.01
PaRE	7.4E-3 $\pm$ 1E-4	0.286 $\pm$ 4E-3	0.28 $\pm$ 7E-3	6.70 $\pm$ 0.3	11.21 $\pm$ 0.07	27.04 $\pm$ 0.7	7.12 $\pm$ 0.3	0.12 $\pm$ 4E-3	0.99 $\pm$ 0.03	0.55 $\pm$ 0.01
MoNA	6.8E-3	0.280	0.27	6.48	11.13	27.13	7.28	0.121	0.99	0.55
RECRAFT	7.2E-3 $\pm$ 7E-5	0.278 $\pm$ 2E-3	0.27 $\pm$ 6E-3	7.30 $\pm$ 0.4	11.11 $\pm$ 0.04	26.41 $\pm$ 0.7	6.60 $\pm$ 0.3	0.11 $\pm$ 3E-3	0.99 $\pm$ 0.02	0.55 $\pm$ 0.01

This evaluation is conducted with respect to the pre-trained source encoder θ , the target encoder ϕ as well as the corresponding target prediction heads $p_\tau(z' | \mathbf{u})$ which are learned via optimizing the above bound using our proposed algorithm in the main text. We will show that the bound is sufficiently tight for the learned target encoder ϕ which consequently demonstrates that the bound in Theorem 7 is a sufficiently good surrogate to optimize for the (inaccessible) target generalization loss.

To achieve this, we use the source’s and target’s *test sets*, which were not used during the pre-training and fine-tuning of the source and target models, to compute the corresponding target loss $\text{err}_\tau(\phi)$ and source loss $\text{err}_s(\theta)$. To compute the bound on the target generalization loss, we calculate the feature alignment (**FA**) using Eq. (14) and approximate the feature-label distortion (**FLD**) using the upper bound in Eq. (21). We then compute the optimal target fitting term **TF** at the learned target feature map ϕ and prediction head $p_\tau(z' | \mathbf{u})$ on the target test set via (1) solving for $\Lambda_{\mathbf{u}}^+(\cdot | \cdot)$ that minimizes the feature-label distortion (**FLD**),

$$\Lambda_{\mathbf{u}}^+(\cdot | \cdot) = \underset{\Lambda_{\mathbf{u}} \in C_{\mathbf{u}}^*}{\text{argmin}} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[\mathbb{H}[\Lambda_{\mathbf{u}}^*(\cdot | z)] \right], \quad (46)$$

with respect to the linear constraint in Eq. (6); and (2) leveraging it to equivalently rewrite the **TF** formulation in Definition 6 as the optimal solution for a convex optimization task with linear constraints:

$$\begin{aligned} \mathbf{TF}(\mathbf{u}) &= \min_{\Lambda_{\mathbf{u}}(\cdot | \cdot)} \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[\mathbb{KL}(\Lambda_{\mathbf{u}}^+(\cdot | z) \| \Lambda_{\mathbf{u}}(\cdot | z)) \right] \\ \text{subject to } p_\tau(z' | \mathbf{u}) &= \mathbb{E}_{D_s^\theta(z|\mathbf{u})} \left[\Lambda_{\mathbf{u}}(z' | z) \right] \text{ for each } z' \in \mathcal{Z}'. \end{aligned} \quad (47)$$

The above is a direct consequence of the formulation in Definition 6 when we fix a particular choice of the target’s feature map ϕ and its corresponding prediction head $p_\tau(z' | \mathbf{u})$.

Importantly, we note that such formulation assumes knowledge of ϕ and $p_\tau(z' | \mathbf{u})$ as well as the test set and therefore cannot be used as an alternative to our proposed algorithm for learning ϕ and $p_\tau(z' | \mathbf{u})$. Instead, this formulation is an effective probing tool for post-training inspection/evaluation of the tightness of the theoretical bound in Theorem 7 as it admits a closed-form solution for $\Lambda_{\mathbf{u}}$ in terms of the (learned) target feature map ϕ , prediction head $p_\tau(z' | \mathbf{u})$, and the corresponding target test data’s induced conditional distribution $D_\tau^\phi(z' | \mathbf{u})$:

$$\Lambda_{\mathbf{u}}(z' | z) = \Lambda_{\mathbf{u}}^+(z' | z) \cdot p_\tau(z' | \mathbf{u}) / D_\tau^\phi(z' | \mathbf{u}) \quad \text{and} \quad \mathbf{TF}(\mathbf{u}) = \mathbb{KL}(D_\tau^\phi(z' | \mathbf{u}) \| p_\tau(z' | \mathbf{u})). \quad (48)$$

Remark. Eq. (48) also provides direct support for the claim in Section 3.2 that minimizing the target loss also minimizes the target fitting error (**TF**). To elaborate, given a feature map ϕ and a specified target task, the oracle conditional distribution $D_\tau^\phi(z' | \mathbf{u})$ is fixed. During training phase in Section 3.2, the target predictor $p_\tau(z' | \mathbf{u})$ is trained to minimize the target loss, which, in turn, reduces the Kullback-Leibler (KL) divergence between $p_\tau(z' | \mathbf{u})$ and $D_\tau^\phi(z' | \mathbf{u})$. This corresponds exactly to Eq. (48).

Using the above calculations, we now have access to both the target’s generalization loss and its upper-bound based on the source’s generalization loss and other key quantities of cross-modal fine-tuning such as feature alignment(**FA**), feature-label distortion(**FLD**) and target fitting(**TF**). The gap between the target’s generalization loss and its upper-bound is visualized in Fig. 5 which shows that our error bound is sufficiently tight with small different gap observed consistently across a variety of target datasets. The averaged gap over all these target tasks is less than 29%.

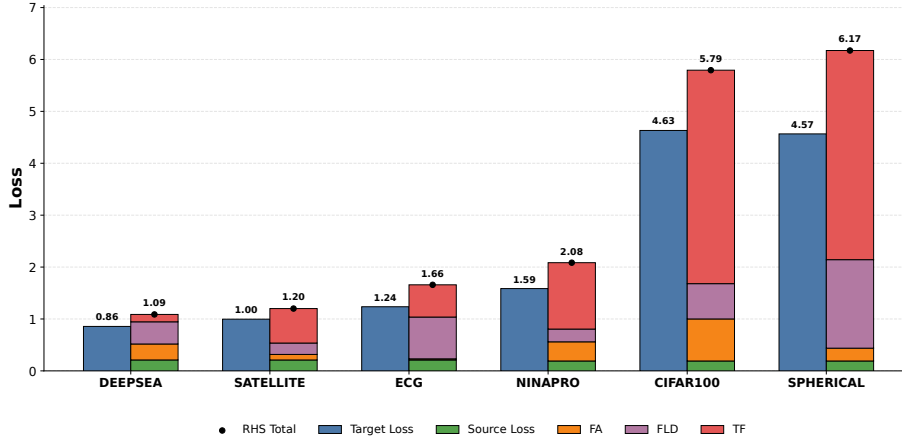


Figure 5: Bar charts illustrating the gap between the target’s generalization loss and its upper-bound established in Theorem 7. For each target task, there are two bars representing the target’s generalization loss and its upper-bound. The bar representing the upper-bound is further demarcated into the source’s generalization loss (fine-tuning overhead), feature alignment(**FA**), feature Label distortion(**FLD**), and target fitting(**TF**). The bound evaluation is conducted at the learned target’s feature map ϕ and its corresponding prediction head $p_\tau(z' | \mathbf{u})$.

C Related Works

This section summarizes the existing literature on domain adaptation (Appendix C.1), in-modal fine-tuning (Appendix C.2), and cross-modal fine-tuning (Appendix C.3).

C.1 Domain Adaptation

Domain adaptation is a form of transductive transfer learning in which the inputs and outputs of the source and target tasks are identical, but their corresponding data distributions are different (Pan and Yang, 2010; Wang and Deng, 2018). Its main focus is on transferring knowledge within the same task across different data distributions (i.e., domains) assuming shared input/output spaces. It also assumes that target inputs are accessible during training (often in the absence of target labels) to enable distribution alignment in a transductive setting. Furthermore, prior to the advent of foundation models (FMs), algorithmic designs for domain adaptation were tailored to specific task (i.e., input/output spaces) and source model’s architecture, limiting their applicability to other scenarios. In contrast, fine-tuning in the FM paradigm enables model-agnostic adaptation procedures that operate without assuming specific forms of input/output or requiring access to target data during training. While some recent efforts in domain adaptation have relaxed assumptions on shared input/output spaces (Hoang et al., 2020; Lam et al., 2021; Hoang and Hoang, 2024; Yin et al., 2024), they still require source and target inputs to originate from the same modality (e.g., images, text), which restricts their use in more general cross-modal transfer settings. Other recent approaches have considered text-to-image transfer but require access to the (unlabeled) target data during training, are restricted to the same task, use a specific model architecture catering to text and image modalities (Yao et al., 2019; Li et al., 2020). This is different from the broader

scheme of cross-modal fine-tuning which aims to enable transfer across different tasks and unseen modalities via model-agnostic procedures.

C.2 In-Modal Fine-Tuning of FMs

The recent advent of large pre-trained or foundation models (FMs) has enabled flexible knowledge transfer to a wide range of downstream tasks under a unified, task-agnostic paradigm. This is commonly known as fine-tuning, which is model-agnostic and does not require shared input or output spaces. Such transfer is feasible because these FMs are trained on broad, diverse corpora (e.g., text, image, or audio) that often overlap semantically or structurally with the content of downstream datasets for in-modal tasks. As a result, fine-tuning has been successfully applied across domains such as vision (Kirillov et al., 2023; Liu et al., 2021b), video understanding (Bertasius et al., 2021), language (Zheng et al., 2021; Yang et al., 2022; Ma et al., 2023), and speech (Radford et al., 2022; Li et al., 2021a). There are also multi-modal FMs (Radford et al., 2021; Alayrac et al., 2022; Kim et al., 2021) which were pre-trained to learn the embeddings of multiple modalities together but existing approaches to facilitate knowledge transfer from these models are still restricted to within the pre-trained modalities. Extending fine-tuning to cross-modal scenarios with unseen data and downstream tasks (e.g., fine-tuning vision/language FMs to solve complex physics problems) are however less explored as discussed next.

C.3 Cross-Modal Fine-Tuning of FMs

Early attempts in cross-modal fine-tuning have focused on transferring language models to other modalities such as vision (Kiela et al., 2019; Tan and Bansal, 2019; Gu et al., 2022), DNA/protein sequences (Nguyen et al., 2023; Lin et al., 2023; Jumper et al., 2021). These provide initial evidence of the cross-modal transferability of FMs but their designs are nonetheless hand-tailored to specific target tasks and modalities rather than for general-purpose (Shen et al., 2023). More recently, a few general-purpose cross-modal fine-tuning frameworks have emerged with remarkable successes in finetuning existing vision (Liu et al., 2021b) or language (Liu et al., 2019) FMs to solve a broad and diverse set of unseen tasks and data modalities (Shen et al., 2023; Cai et al., 2024; Ma et al., 2024). In particular, ORCA (Shen et al., 2023) aligns source and target representation distributions using optimal transport before fine-tuning the entire network. PARE (Cai et al., 2024) alternatively introduces a gating mechanism to integrate source and target features during fine-tuning. MoNA (Ma et al., 2024) addresses modality representation alignment via a bi-level optimization framework: the inner loop learns the optimal target predictor for a given embedder, while the outer loop updates the embedder to align the source feature-label semantics with the combined representation and prediction under this predictor. This help reduces the misalignment between source’s and target’s feature-label structures while fitting to the target data. However, as noted in Section 1, these approaches largely adopt heuristic methods combining distributional feature alignment and target fitting to facilitate cross-modal knowledge transfer. Overall, these approaches lack a principled means to assess the impact of their heuristic strategies on the generalized target performance.

D Technical Background on Wasserstein Distance

As the core technical block of our theoretical analysis is built on the Wasserstein distance, we provide a succinct definition of this distance (Appendix D.1) and its computation (Appendix D.2) below.

D.1 Wasserstein Distance

The Wasserstein p -distance is a fundamental metric in the theory of optimal transport. It is used to measure the distance between probability measures defined on a metric space. Intuitively, it is the cost of transporting mass from one distribution to another with respect to a cost metric $\delta(\mathbf{u}_1, \mathbf{u}_2)$ measuring the distance between two locations. This concept is widely applied in fields such as probability theory, machine learning, and partial differential equations (Villani, 2008). Let $(\mathcal{U}, \delta)$ be a metric space that is a Polish space. For $p \in [1, +\infty]$, the Wasserstein- p distance between two probability measures μ and ν on space $\mathcal{U}$ is:

$$W_\delta^p(\mu, \nu) = \inf_{\pi \in \Pi(\mu, \nu)} \left(\mathbb{E}_{(\mathbf{u}_1, \mathbf{u}_2) \sim \pi} [\delta(\mathbf{u}_1, \mathbf{u}_2)^p] \right)^{\frac{1}{p}} \quad (49)$$

For example, we denote the Wasserstein-1 distance between two probability measures μ and ν on $(\mathcal{U}, \delta)$ as:

$$W_\delta(\mu, \nu) = \inf_{\pi \in \Pi(\mu, \nu)} \left(\mathbb{E}_{(\mathbf{u}_1, \mathbf{u}_2) \sim \pi} [\delta(\mathbf{u}_1, \mathbf{u}_2)] \right) \quad (50)$$

$$= \inf_{\pi \in \Pi(\mu, \nu)} \int_{\mathcal{U}} \int_{\mathcal{U}} \delta(\mathbf{u}_1, \mathbf{u}_2) \pi(\mathbf{u}_1, \mathbf{u}_2) d\mathbf{u}_1 d\mathbf{u}_2. \quad (51)$$

A key property of the Wasserstein-1 distance is the Kantorovich duality theorem, which provides a dual formulation of the distance in terms of Lipschitz functions. This duality is particularly useful in applications ranging from machine learning to functional analysis (Villani, 2008) and is stated below. For any $\tau_\delta > 0$, let Lip_{τ_δ} denote the set of functions $g : \mathcal{U} \rightarrow \mathbb{R}$ such that

$$|g(\mathbf{u}_1) - g(\mathbf{u}_2)| \leq \tau_\delta \delta(\mathbf{u}_1, \mathbf{u}_2)$$

for all $\mathbf{u}_1, \mathbf{u}_2 \in \mathcal{U}$. Then, the Wasserstein-1 distance is given by

$$W_\delta(\mu, \nu) = \frac{1}{\tau_\delta} \sup_{f \in \text{Lip}_{\tau_\delta}} \left(\mathbb{E}_{\mathbf{u}_1 \sim \mu} [g(\mathbf{u}_1)] - \mathbb{E}_{\mathbf{u}_2 \sim \nu} [g(\mathbf{u}_2)] \right). \quad (52)$$

D.2 Computation of Wasserstein Distance

Computing the Wasserstein distance is achieved via solving a linear program with a computational complexity of $O(n^3)$ (Peyré and Cuturi, 2020). This is however impractical for large datasets. To address this, we adopt entropic regularization to introduce a penalty term $\epsilon H(\pi) = -\epsilon \mathbb{E}_\pi [\log(\pi)]$ with $\epsilon > 0$ which leads to the following augmented optimization:

$$\pi_\epsilon = \arg \min_{\pi \in \Pi(\mu, \nu)} \left(\mathbb{E}_{(\mathbf{u}_1, \mathbf{u}_2) \sim \pi} [\delta(\mathbf{u}_1, \mathbf{u}_2)] + \epsilon H(\pi) \right). \quad (53)$$

This regularized problem is solved efficiently using the Sinkhorn algorithm, which iteratively scales the transport plan $\pi_\epsilon = \text{diag}(u)K\text{diag}(v)$, where u, v are scaling vectors updated to satisfy the marginal constraints.

The algorithm has an overall complexity of $O(n^2)$ which offers significant speedup for large-scale problems (Peyré and Cuturi, 2020). For a practical implementation, we utilize the Sinkhorn algorithm as provided by the Python Optimal Transport (POT) library and use its default regularization parameter $\epsilon = 0.1$ to ensure a balance between computational efficiency and accuracy.

E Additional Experimental Results and Ablation Studies

This section provides additional experimental results to supplement the main text results. These results include: (i) the performance of RECRAFT compared to specialized physic-informed methods such as PINN (Raissi et al., 2019) and FNO (Li et al., 2021b), as well as a baseline U-Net method (Ronneberger et al., 2015) on PDEBench as shown in Table 3; (ii) experiment results on NAS-Bench-360 (Tu et al., 2022) with additional details on error bars for each task, as shown in Table 4; (iii) experiment results on PDEBench with additional details on error bars for each task, as shown in Table 5 and Table 6.

We also provide the ablations across several tasks on NAS-Bench-360 (Tu et al., 2022) isolating contributions: NFT vs. FA-only vs. RECRAFT, as shown in Table 7. RECRAFT achieves the best performance across all tasks and highlights the contribution of feature-label distortion (FLD) loss.

To assess the sensitivity of ω , which ensures the Lipschitz continuity of the loss function ℓ_s on the source dataset and balances the minimization of feature alignment (**FA**) and feature-label distortion (**FLD**), we evaluate the model’s performance across different target datasets using varying values of ω . Fig. 6 illustrates the sensitivity of ω across various target datasets. The results suggest that setting ω within the range of approximately 0.3 to 0.5 achieves an optimal balance between **FA** and **FLD** during optimization.

F Practical Source Re-calibration to Ensure Lipschitz Constraint with ℓ_2 Metric

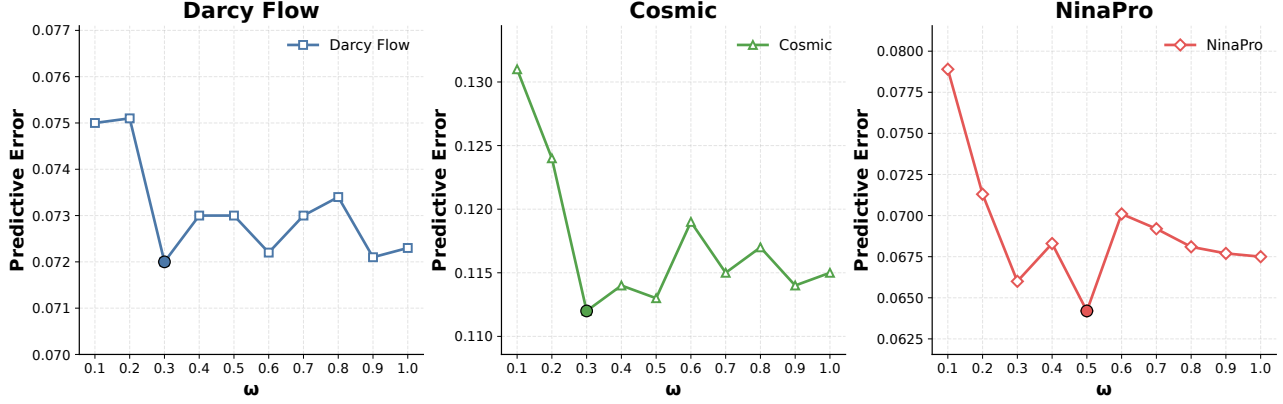


Figure 6: Predictive error ($\downarrow$) of Darcy Flow, Cosmic and Ninapro across various values of ω , which achieves the balance between minimizing Feature Alignment(**FA**) and Feature Label Distortion(**FLD**).

Table 5: Predictive error ($\downarrow$) (with standard deviation) achieved by our proposed method RECRAFT and other specialized physic-informed baselines across multiple PDE tasks. U-Net results for Naiver-Stokes and Darcy Flow are missing because the benchmark paper (Takamoto et al., 2022) does not evaluate them (due to memory issues).

Model	Darcy nRMSE	Advection nRMSE	Burgers nRMSE	Diff.Sorp nRMSE	S.Water nRMSE	Diff.Reac(2D) nRMSE	Diff.Reac(1D) nRMSE	Navier-Stokes nRMSE
PINN	$0.180 \pm 2E-3$	0.67 ± 0.03	0.36 ± 0.03	0.15 ± 0.03	$0.085 \pm 3E-3$	0.840 ± 0.01	0.840 ± 0.01	$0.720 \pm 5E-3$
FNO	$0.220 \pm 3E-3$	$0.011 \pm 1E-3$	$0.0031 \pm 5E-5$	$1.8E-3 \pm 4E-5$	$4.4E-3 \pm 3E-5$	$0.120 \pm 4E-4$	$1.4E-3 \pm 1E-4$	$0.068 \pm 2E-3$
U-Net	-	1.1 ± 0.21	0.99 ± 0.02	0.22 ± 0.02	$0.017 \pm 2E-3$	$1.6 \pm 7E-3$	$0.08 \pm 4E-3$	-
RECRAFT	$0.079 \pm 1E-3$	$0.0078 \pm 4E-4$	$0.0108 \pm 3E-4$	$1.6E-3 \pm 3E-4$	$5.4E-3 \pm 5E-5$	$0.817 \pm 3E-3$	$2.8E-3 \pm 2E-4$	$0.050 \pm 3E-3$

This section provides a concrete example of how the source’s prediction map can be recalibrated to ensure Lipschitz constraint with a practical choice of δ . Note that the theoretical bound in Theorem 7 holds with any metric, we can then choose δ to be a Euclidean distance, i.e., $\delta \equiv \ell_2$. Given the choice, we can now aim to minimize $\tau_\delta = \omega$ via recalibrating the source’s prediction map. This is achieved via finding a smallest value of ω for which the source’s prediction map can be refitted so that $\ell_s(\mathbf{u}) := -\mathbb{E}_{D_s^q(z|\mathbf{u})}[\log(p_s(z|\mathbf{u}))]$ is $O(\omega)$ -Lipschitz with respect to the Euclidean metric $\delta \equiv \ell_2$ while preserving performance on the proxy dataset. To do this, we view ω as a hyper-parameter and run ablation experiments on the proxy dataset with different choices of ω over a pre-defined range $[0.1, 1.0]$ to determine its optimal choice. For each ω , we recalibrate the source’s prediction map via

$$\gamma(\omega) = \operatorname{argmin}_{\gamma} \mathbb{E}_{\mathbf{x} \sim P_s} \max \left(0, \|\nabla_{\mathbf{u}} \ell_s(\theta(\mathbf{x}); \gamma)\|_2 - \omega \right)$$

with γ denotes the last layer of $p_s(z|\mathbf{u})$. Once the re-calibrated $\gamma(\omega)$ has been computed, we re-evaluate the source performance to determine whether $\gamma(\omega)$ preserves performance and ω can be further reduced.

Our experiments show that smaller values of ω indeed lead to larger source loss values as they impose more constraints on the prediction map. Based on these ablation results, we selected $\omega = 0.3$ for 2D tasks; and $\omega = 0.5$ for 1D tasks. These are the minimum values of ω for which the re-calibrated prediction map preserves the source’s model on the proxy dataset. To visualize this, we present Fig. 7 which plots the re-calibrated (source) model’s performance on the proxy CIFAR-10 dataset versus Lipschitz threshold $\tau_\delta \triangleq \omega \in [0.1, 1.0]$ (for 2D tasks).

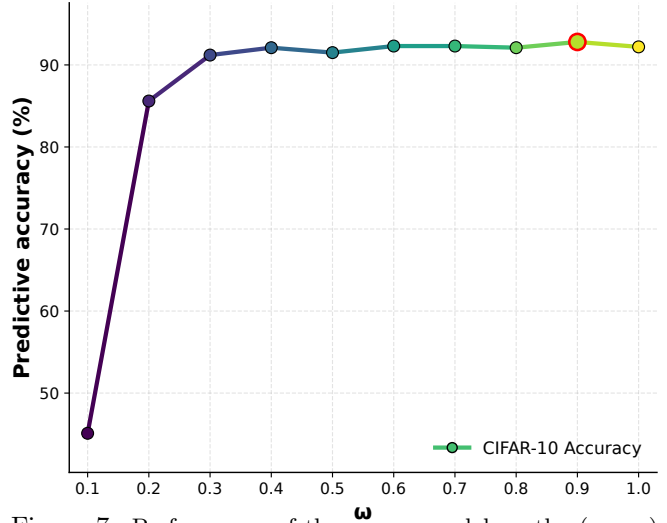


Figure 7: Performance of the source model on the (proxy) CIFAR-10 dataset with re-calibrated prediction head to satisfy the Lipschitz constraint in Definition 4 with constant $\tau_\delta = \omega$ where ω varies in $[0.1, 1.0]$.

Table 6: Predictive error ($\downarrow$) (with standard deviation) achieved by our proposed method RECRAFT and other state-of-the-art (SOTA) cross-modal fine-tuning baselines across multiple PDE tasks in PDEBench (Takamoto et al., 2022). The reported results of MoNA (Ma et al., 2024) are sourced from the corresponding paper due to its unavailable code.

Model	Darcy nRMSE	Advection nRMSE	Burgers nRMSE	Diff.Sorp nRMSE	S.Water nRMSE	Diff.Reac(2D) nRMSE	Diff.Reac(1D) nRMSE	Navier-Stokes nRMSE
NFT	$0.085 \pm 4E-3$	$0.0140 \pm 2E-3$	$0.0130 \pm 4E-4$	$3.1E-3 \pm 7E-5$	$6.1E-3 \pm 4E-5$	$0.830 \pm 4E-3$	$9.2E-3 \pm 2E-3$	$0.863 \pm 2E-2$
ORCA	$0.081 \pm 1E-3$	$0.0098 \pm 2E-4$	$0.0120 \pm 4E-4$	$1.8E-3 \pm 2E-4$	$6.0E-3 \pm 5E-5$	$0.820 \pm 3E-3$	$3.2E-3 \pm 2E-4$	$0.066 \pm 2E-3$
PARE	$0.081 \pm 1E-3$	$0.0032 \pm 4E-4$	$0.0114 \pm 3E-4$	$1.9E-3 \pm 3E-4$	$5.9E-3 \pm 5E-5$	$0.820 \pm 5E-3$	$2.9E-3 \pm 3E-4$	$0.068 \pm 3E-3$
MoNA	0.079	0.0088	0.0114	$1.6E-3$	5.7E-3	0.818	$2.8E-3$	0.054
RECRAFT	$0.079 \pm 1E-3$	$0.0078 \pm 4E-4$	$0.0108 \pm 3E-4$	$1.6E-3 \pm 3E-4$	$5.4E-3 \pm 5E-5$	$0.817 \pm 3E-3$	$2.8E-3 \pm 2E-4$	$0.050 \pm 3E-3$

Table 7: Prediction errors ($\downarrow$) incurred by NFT(naive fine-tuning),FA(only optimize feature alignment) and RECRAFT (with standard deviation) across 10 diverse tasks on NAS-Bench-360. Our method RECRAFT achieves best performance all tasks, highlight the effectiveness of minimizing FLD.

Model	Darcy Relative ℓ_2	DeepSEA 1- AUROC	ECG 1-F ₁ score	CIFAR100 0-1 error (%)	Satellite 0-1 error (%)	Spherical 0-1 error(%)	Ninapro 0-1 error (%)	Cosmic 1- AUROC	Psicov MAE _s	FSD50K 1-mAP
NFT	$7.4E-3 \pm 1E-4$	$0.490 \pm 2E-2$	0.44 ± 0.03	9.74 ± 1.21	13.82 ± 0.24	55.76 ± 2.3	8.35 ± 0.4	$0.17 \pm 2E-2$	1.92 ± 0.06	0.63 ± 0.01
FA	$7.3E-3 \pm 8E-5$	$0.290 \pm 2E-3$	$0.29 \pm 6E-3$	7.80 ± 0.41	11.61 ± 0.20	29.85 ± 0.8	7.73 ± 0.4	$0.16 \pm 5E-3$	1.92 ± 0.04	0.57 ± 0.01
Ours	$7.2E-3 \pm 7E-5$	$0.278 \pm 2E-3$	$0.27 \pm 6E-3$	7.30 ± 0.4	11.11 ± 0.04	26.41 ± 0.7	6.60 ± 0.3	$0.11 \pm 3E-3$	0.99 ± 0.02	0.55 ± 0.01

G RECRAFT algorithm

For better clarity, this section provides the pseudocode for our algorithm RECRAFT previously described in Section 3. To summarize, RECRAFT adopts a two-stage approach that decomposes the bound minimization into (1) finding a feature map ϕ that minimizes the semantic gap $\mathbf{FA}(\phi, \theta) + \mathbb{E}_{D_{\tau}^{\phi}(\mathbf{u})}[\mathbf{FLD}(\mathbf{u})]$ between the source and target tasks, thus maximizing transferability (Section 3.1); and (2) learning a predictor $p_{\tau}(z' | \mathbf{u})$ based on the learned ϕ (Section 3.2) via minimizing $\mathbb{E}_{D_{\tau}^{\phi}(\mathbf{u})}[\mathbf{TF}(\mathbf{u})]$. A pseudocode of RECRAFT is detailed in Algorithm 1.

Algorithm 1 RECRAFT Algorithm

- 1: **Input:** Pre-trained model $M_s \triangleq (\theta, p_s(z | \theta(\mathbf{x})))$, proxy dataset $(\mathbf{X}^s, \mathbf{z}^s) \sim D_s(\mathbf{x}, z)$, target dataset $(\mathbf{X}^{\tau}, \mathbf{z}^{\tau}) \sim D_{\tau}(\mathbf{x}', z')$, number of epochs n_0, n_1, n_2 .
 - 2: **Output:** The target model $M_{\tau} \triangleq (\phi, p_{\tau}(z' | \phi(\mathbf{x}')))$.
 - 3: **Stage 1: Learning Feature Map (Section 3.1)**
 - 4: Initialize the target embedder ϕ .
 - 5: **for** $epoch = 1$ to n_1 **do**
 - 6: Compute $L_{\mathbf{FA}}(\phi)$ using Eq. 14.
 - 7: Update ϕ by minimizing $L_{\mathbf{FA}}(\phi)$.
 - 8: **end for**
 - 9: **for** $epoch = 1$ to n_2 **do**
 - 10: Compute $L_{\mathbf{FLD}}(\phi)$ using Eq. 21.
 - 11: Update ϕ by minimizing $L_{\mathbf{FLD}}(\phi)$.
 - 12: **end for**
 - 13: **Stage 2: Learning Target Predictor (Section 3.2)**
 - 14: Frozen ϕ , initialize $p_{\tau}(z' | \mathbf{u})$ from $p_s(z | \mathbf{u})$.
 - 15: **for** $epoch = 1$ to n_0 **do**
 - 16: Compute $L_{\mathbf{TF}}$ using Eq. 24.
 - 17: Update p_{τ} by minimizing $L_{\mathbf{TF}}$.
 - 18: **end for**
 - 19: **Return** $(\phi, p_{\tau}(z' | \mathbf{u}))$
-

H Hyperparameters & Implementation Details

To ensure fair comparisons, we adopt the same hyperparameters as those used in **ORCA** (Shen et al., 2023) for model fine-tuning. These specific parameter settings are shown in Table 8 for Nas-Bench-360 (Tu et al., 2022)

and Table 9 for PDEBench (Takamoto et al., 2022). For detailed implementation, we construct the target model $M_\tau \triangleq (\phi, p_\tau(z' | \phi(\mathbf{x}')))$ by adapting a pre-trained foundation model $M_s \triangleq (\theta, p_s(z | \theta(\mathbf{x})))$ which is selected to be a Swin Transformer (Liu et al., 2021b) for 2D tasks; and a RoBERTa (Liu et al., 2019) for 1D tasks. To compute the FLD for a regression task, we discretize the continuous label space into 10 distinct classes, consistent with the hyperparameter settings used in ORCA (Shen et al., 2023).

All experiments are run on an NVIDIA RTX 2080 GPU and results are averaged over 5 independent runs. Our code repository can be found at <https://github.com/khiembk/RECRAFT>.

Table 8: Hyperparameter configurations used for the 10 tasks on NAS-Bench-360 (Tu et al., 2022). *FA epoch* and *FLD epoch* denote the number of training epochs used to minimize the feature alignment (**FA**) and feature-label distortion (**FLD**). The optimizers used in our experiments include SGD, Adam (Kingma and Ba, 2017) and AdamW (Loshchilov et al., 2017).

Hyperparameter	CIFAR100	Spherical	NinaPro	FSD50K	Darcy Flow	PSICOV	Cosmic	ECG	Satellite	DeepSEA
Backbone	Swin	Swin	Swin	Swin	Swin	RoBERTa	Swin	RoBERTa	RoBERTa	RoBERTa
Batch size	32	32	32	32	4	1	4	4	16	16
Epoch	60	60	60	100	100	10	60	15	60	13
Accumulation	32	4	1	1	1	32	1	16	4	1
Optimizer	SGD	AdamW	Adam	Adam	AdamW	Adam	AdamW	SGD	AdamW	Adam
Learning rate	1.00E-04	1.00E-04	1.00E-04	1.00E-04	1.00E-03	5.00E-06	1.00E-03	1.00E-06	3.00E-05	1.00E-05
Weight decay	1.00E-03	1.00E-01	1.00E-05	5.00E-05	5.00E-03	1.00E-05	0.00E+00	1.00E-01	3.00E-06	0.00E+00
Label discretization	-	-	-	-	10	10	10	10	-	-
ω	0.3	0.3	0.3	0.3	0.3	0.5	0.3	0.5	0.5	0.5
FA epoch	60	60	60	60	60	60	60	60	60	60
FLD epoch	4	4	4	4	4	4	4	2	4	2

Table 9: Hyperparameter configurations used for the 8 tasks on PDEBench (Takamoto et al., 2022). *FA epoch* and *FLD epoch* denote the number of training epochs used to minimize the feature alignment (**FA**) and feature-label distortion (**FLD**). The optimizers used in our experiments include SGD, Adam (Kingma and Ba, 2017) and AdamW (Loshchilov et al., 2017).

Hyperparameter	Advection	Burgers	RD1D	Diff-Sor	Navier-Stokes	Darcy-Flow	Shallow-Water	RD2D
Backbone	RoBERTa	RoBERTa	RoBERTa	Swin	RoBERTa	Swin	Swin	Swin
Batch size	4	4	4	4	4	4	4	4
Epoch	200	200	200	200	200	100	200	200
Accumulation	1	1	1	1	1	1	1	1
Optimizer	Adam	Adam	SGD	AdamW	AdamW	AdamW	AdamW	Adam
Learning rate	1.00E-04	1.00E-05	1.00E-03	1.00E-04	1.00E-04	1.00E-04	1.00E-04	1.00E-04
Weight decay	1.00E-05	1.00E-05	1.00E-05	0	1.00E-03	1.00E-05	0	1.00E-03
Label discretization	10	10	10	10	10	10	10	10
ω	0.5	0.5	0.5	0.3	0.5	0.3	0.3	0.3
FA epoch	60	60	60	60	60	60	60	60
FLD epoch	4	4	4	4	4	4	4	4

I Complexity Analysis

We provide the complexity analysis of our proposed method as follows:

Stage 1: Learning Feature Map: Computing the loss $L_{\mathbf{FA}}(\phi)$ in Eq. (14) incurs $O(f_\phi * |D_\tau| + f_\theta * |D_s| + |W_\delta|)$ processing cost where f_ϕ and f_θ are the forward costs of the source and the target embedder, respectively; $O(|W_\delta|)$ is the cost of computing the Wasserstein distance W_δ . Therefore, learning the target embedder within n_1 epochs with loss $L_{\mathbf{FA}}(\phi)$ requires $O(f_\theta * |D_s| + n_1(f_\phi * |D_\tau| + |W_\delta| + b_\phi))$ where b_ϕ is the cost of backpropagation of the target embedder. Likewise, updating the target embedder within n_2 epochs with loss $L_{\mathbf{FLD}}(\phi)$ in Eq. (21) incurs $O(n_2(f_{p_s} * |D_\tau| + t_{FLD} + b_{p_s, \phi}))$ where f_{p_s} is the cost of attaining $z \sim p_s(z | \phi(\mathbf{x}'))$, t_{FLD} is the cost of computing the conditional entropy in Eq. (21), and $b_{p_s, \phi}$ is the cost of backpropagating through the frozen source predictor and target embedder.

Stage 2: Learning Target Predictor: Updating the target predictor with the loss $L_{\mathbf{TF}}$ in n_0 epochs requires $O(n_0(f_{p_\tau} * |D_\tau| + |D_\tau| + b_{p_\tau}))$ where f_{p_τ} and b_{p_τ} are forward and backpropagation cost of the target predictor.

[illegible]

Table 11: Data statistics of the 10 tasks in NAS-Bench-360 (Tu et al., 2022).

	CIFAR100	Spherical	NinaPro	FSD50K	Darcy Flow	PSICOV	Cosmic	ECG	Satellite	DeepSEA
# Training data	60K	60K	43956	51K	1.1K	3606	5250	330K	1M	250K
Input shape	2D	2D	2D	2D	2D	1D	2D	1D	1D	1D
Output type	Point	Point	Point	Point	Dense	Dense	Dense	Point	Point	Point
# Classes	100	100	18	200	–	–	–	4	24	36
Loss	CE	CE	LpLoss	MSELoss	BCE	FocalLoss	BCE	CE	CE	BCE
Expert Network	DenseNet-BC	S2CN	Attention Model	VGG	FNODE	EPCON	deepCR-mask	ResNet-1D	ROCKET	DeepSEA

PDEBench (Takamoto et al., 2022) comprises multiple scientific datasets with simulated data from a wide variety of partial differential equations (PDEs) in physics. These include the 1D Advection equation, modeling linear advection with a constant speed parameter β ; the 1D Burgers’ equation, capturing non-linear fluid dynamics with a constant diffusion coefficient ν ; the 1D Diffusion-Reaction equation, combining diffusion and a source term governed by parameters ν and ρ ; and the 1D Diffusion-Sorption equation, representing diffusion retarded by sorption, applicable to real-world scenarios. Additionally, the 1D Navier-Stokes equation describes compressible fluid dynamics with shear and bulk viscosities η and ζ . In two dimensions, the Darcy-Flow equation models steady-state flow over a unit square, scaled by a constant force term β ; the Shallow-Water equations, derived from Navier-Stokes, address free-surface flow problems; and the 2D Diffusion-Reaction equation extends its 1D counterpart with two non-linearly coupled variables, presenting a complex challenge with significant real-world applications. Table 10 presents the details of each task, with a more comprehensive description available in the original paper (Takamoto et al., 2022).

K Boarder Impacts

Our research provides a new analytical lens that opens several promising research directions:

(i) Knowledge Distillation (KD). The bound naturally extends to KD by decomposing the teacher–student gap into feature alignment (FA) and softened-label distribution mismatch (FLD). This perspective suggests a new class of KD objectives that jointly minimize both components, potentially yielding more effective and robust distillation compared to standard temperature-scaled KL or feature-mimicry losses.

(ii) Retrieval-Augmented Generation (RAG) and Multimodal Retrieval. The same decomposition can guide the alignment of retrieved documents or heterogeneous data modalities (e.g., audio, video, or scientific figures) within a shared representation space. This offers a theoretically grounded alternative to CLIP-style contrastive learning, with the potential for improved computational efficiency and robustness under label shift.

(iii) Scaling to Foundation Models and LLMs. Our framework highlights that explicitly accounting for FLD, beyond standard feature alignment, is critical for effective transfer. This insight provides a principled foundation for efficient cross-modal fine-tuning in large-scale models. Preliminary results on pre-trained Swin Transformer (90M parameters) and RoBERTa (88M+ parameters) demonstrate the feasibility of this approach and suggest a clear path toward scaling to modern foundation models.