

UNDERSTANDING THE CAPABILITIES AND LIMITATIONS OF WEAK-TO-STRONG GENERALIZATION

Wei Yao^{1*}, Wenkai Yang^{1*}, Ziqiao Wang², Yankai Lin¹, Yong Liu^{1†}

¹ Gaoling School of Artificial Intelligence, Renmin University of China,

² School of Computer Science and Technology, Tongji University,
{wei.yao, liuyonggsai}@ruc.edu.cn

ABSTRACT

Weak-to-strong generalization, where weakly supervised strong models outperform their weaker teachers, offers a promising approach to aligning superhuman models with human values. To deepen the understanding of this approach, we provide theoretical insights into its capabilities and limitations. First, in the classification setting, we establish upper and lower generalization error bounds for the strong model, identifying the primary limitations as stemming from the weak model’s generalization error and the optimization objective itself. Additionally, we derive lower and upper bounds on the calibration error of the strong model. These theoretical bounds reveal two critical insights: (1) the weak model should demonstrate strong generalization performance and maintain well-calibrated predictions, and (2) the strong model’s training process must strike a careful balance, as excessive optimization could undermine its generalization by over-relying on the weak supervision. Finally, in the regression setting, we extend the work of Charikar et al. (2024) to a loss function based on KL divergence, offering guarantees that the strong student can outperform its weak teacher by at least the magnitude of their disagreement. The theory is validated through sufficient experiments.

1 INTRODUCTION

Recently, human supervision (Ouyang et al., 2022; Bai et al., 2022a) plays a crucial role in building both effective and safe artificial intelligence systems (Achiam et al., 2023; Touvron et al., 2023). However, as future superhuman models exhibit increasingly complex behaviors, reliable human oversight becomes increasingly challenging (OpenAI, 2024).

To tackle this issue, the Weak-To-Strong Generalization (WTSG) paradigm (Burns et al., 2023) is proposed. It finds that, strong pre-trained language models, when fine-tuned using labels produced by weaker models, consistently achieve better performance than their weak supervisors. This intriguing phenomenon has not only driven the development of diverse alignment algorithms (Zhu et al., 2025; Liu & Alahi, 2024), but also inspired efforts (Pawelczyk et al., 2024; Yang et al., 2024; Guo et al., 2024) to extend the concept to other tasks. However, despite its empirical success, the theoretical foundations of WTSG remain underdeveloped. Although several elegant theoretical frameworks (Lang et al., 2024; Somerstep et al., 2024; Wu & Sahai, 2025; Charikar et al., 2024) are proposed, a universal framework is still lacking to address fundamental questions, such as: *What is the optimal generalization performance a strong model can achieve after WTSG? Besides generalization, what other factors are influenced by WTSG?*

To answer these questions, we provide a comprehensive theoretical analysis of WTSG, shedding lights on its capabilities and limitations. Firstly, in classification tasks, by assuming that the output of the softmax module lies in the interval $(0, 1]$, our analysis of lower and upper generalization bounds under KL divergence loss reveals that the strong model’s performance is fundamentally determined by two key factors: (1) the disagreement between strong and weak models, which serves as the minimization objective in WTSG, and (2) the weak model’s performance. These findings suggest that (1) achieving the minimal optimization objective in WTSG limits the strong model’s ability to

* Equal contribution † Corresponding author

significantly outperform its weak supervisor, and (2) selecting a stronger weak model can enhance the performance of the strong model. Secondly, we investigate how strong model’s calibration (Guo et al., 2017; Kumar et al., 2019)—the property that a model’s predicted confidence aligns with its actual accuracy—is affected in the WTSG framework. Our theoretical bounds reveal that the calibration of the strong model depends on both the calibration of the weak model and the disagreement between the two models. They highlight the importance of avoiding a poorly-calibrated weak model and an overfitted strong model. The above theory is validated using GPT-2 series (Radford et al., 2019) and Pythia series (Biderman et al., 2023).

In addition to classification setting, we also consider the regression problem. In particular, we build on the work of Charikar et al. (2024) by extending their analysis of squared loss to output distribution divergence, a measure of the difference between two models’ output distributions. In this setting, the model outputs are normalized to form valid probability distributions over all input data, and the output distribution divergence between two models is defined as the KL divergence of their respective output distributions. We recover the findings from Charikar et al. (2024) and show that the strong model’s generalization error is provably smaller than the weak model’s, with the gap no less than the WTSG minimization objective—namely, the strong model’s error on the weak labels. We conduct synthetic experiments to support our theoretical insights.

2 RELATED WORK

In this section, we introduce AI alignment and WTSG. Additional related work including teacher-student learning paradigm, weakly-supervised learning, calibration and information-theoretic analysis is provided in Appendix A.

2.1 AI ALIGNMENT

AI alignment (Ji et al., 2023; Shen et al., 2023) aims to ensure AI systems act in accordance with human values. A popular approach to achieve this goal is fine-tuning models on human-annotated data, such as Reinforcement Learning from Human Feedback (RLHF) (Ouyang et al., 2022; Bai et al., 2022a) and Direct Preference Optimization (DPO) (Rafailov et al., 2024). However, this alignment paradigm faces significant challenges: human oversight becomes insufficient as AI surpasses human capabilities (Kim et al., 2024), and obtaining scalable, high-quality human feedback remains difficult (Casper et al., 2023). These challenges highlight the critical need to align superhuman AI systems (OpenAI, 2024). In contrast to these approaches, our work explores WTSG, which does not rely on extensive human supervision and instead leverages weaker guidance to achieve the alignment goal.

2.2 WEAK-TO-STRONG GENERALIZATION

To explore the effect of weak models to supervise strong models, Burns et al. (2023) first find that strong models supervised by weak models can exhibit better performance on corresponding tasks than their weak supervisors, indicating the possibility of stimulating greater power from super models under weak supervisions. There are also algorithms (Zhu et al., 2025; Agrawal et al., 2024; Sang et al., 2024; Guo & Yang, 2024) and empirical analysis (Yang et al., 2025; Ye et al., 2024) for it. However, to the best of our knowledge, only a limited number of theoretical studies have been conducted on this topic. Lang et al. (2024) analyzes it by introducing theoretical bounds that account for pseudolabel correction and coverage expansion. Somerstep et al. (2024) frame WTSG as a transfer learning problem, revealing limitations of fine-tuning on weak labels. Wu & Sahai (2025) study linear models under a spiked covariance setting and derive asymptotic bounds. Charikar et al. (2024) take a convex-theoretic approach in regression, quantifying performance improvements under squared loss via the misfit error between weak and strong models.

The work most closely related to ours is Charikar et al. (2024), which primarily focuses on squared loss in regression. In contrast, we consider KL divergence-like losses, including KL divergence for classification and output distribution divergence for regression. Furthermore, while they focuses on establishing upper bounds, our study incorporates both upper and lower bounds as well as calibration analysis through experiments on language models.

3 PRELIMINARIES

3.1 CLASSIFICATION AND REGRESSION

We examine two problem settings. In the first case, we consider classification tasks using KL divergence as the loss function. Minimizing this loss is equivalent to minimizing cross-entropy loss, which is widely used in the WTSG literature (Burns et al., 2023). In the second case, we focus on regression tasks, employing the KL divergence between the predictions of two models as the loss function. The model outputs over the entire data domain are normalized to form probability distributions. This approach is an extension of previous result (Charikar et al., 2024) on squared loss, and provides an intuitive framework for understanding WTSG.

Given the data distribution $\mathcal{P}$, data domain $\mathcal{X}$ and output domain $\mathcal{Y}$, let $\mathcal{F} : \mathcal{X} \rightarrow \mathcal{Y}$. Consider the difference $d_{\mathcal{P}}$ and empirical difference $\hat{d}_{\mathcal{P}}$ between two models, where $d_{\mathcal{P}}, \hat{d}_{\mathcal{P}} : \mathcal{F} \times \mathcal{F} \rightarrow \mathbb{R}_0^+$. We define the below two settings:

Setting 1: KL divergence loss. Firstly, we consider a k -classification problem. Given the data domain $\mathcal{X} \subseteq \mathbb{R}^d$ and output domain $\mathcal{Y} \subseteq \mathbb{R}^k$. Consider the model with the softmax module, i.e., $\forall y = (y_1, \dots, y_k)^T \in \mathcal{Y}$, there holds $\sum_{i=1}^k y_i = 1$ and $0 < y_i \leq 1$. Given two models $f, g \in \mathcal{F}$, define $d_{\mathcal{P}}$ and $\hat{d}_{\mathcal{P}}$:

$$d_{\mathcal{P}}(f, g) \triangleq \mathbb{E}_{x \sim \mathcal{P}} [\text{D}_{\text{KL}}(f(x) \| g(x))], \quad (1)$$

$$\hat{d}_{\mathcal{P}}(f, g) \triangleq \frac{1}{n} \sum_{j=1}^n \text{D}_{\text{KL}}(f(x_j) \| g(x_j)), \quad (2)$$

where $\text{D}_{\text{KL}}(f(x) \| g(x)) = \sum_{i=1}^k [f(x)]_i \log \frac{[f(x)]_i}{[g(x)]_i}$ is the KL divergence between predictions, and $[f(x)]_i, [g(x)]_i$ are elements of $f(x), g(x)$.

Setting 2: Output distribution divergence. Secondly, we consider a regression problem. Let the data domain and output domain be $\mathcal{X} \subseteq \mathbb{R}^d$ and $\mathcal{Y} = \{y \in \mathbb{R} | 0 < y \leq 1\}$, respectively. In this setting, the outputs of the model for all input data are probability-normalized to ensure they form valid probability distributions. The difference between two models $f, g \in \mathcal{F}$ is then measured as the KL divergence between their corresponding output distributions:

$$d_{\mathcal{P}}(f, g) \triangleq \int_{\mathcal{X}} f(x) \log \frac{f(x)}{g(x)} dx, \quad (3)$$

$$\hat{d}_{\mathcal{P}}(f, g) \triangleq \sum_{i=1}^n f(x_i) \log \frac{f(x_i)}{g(x_i)}. \quad (4)$$

3.2 WEAK-TO-STRONG GENERALIZATION

In the context of WTSG, we focus on the fine-tuning phase after pre-training. Let $h^* : \mathbb{R}^d \rightarrow \mathbb{R}^{d^*}$ denotes the ground truth representation function, which maps data $x \in \mathcal{X}$ to an ideal, fully enriched representation $h^*(x)$. The target fine-tuning task, composed with the ground truth representation, is denoted as $f^* \circ h^*$, where $f^* : \mathbb{R}^{d^*} \rightarrow \mathcal{Y}$. The *weak model* learns a mapping $f_w \circ h_w$, where the pre-trained representation $h_w : \mathcal{X} \rightarrow \mathbb{R}^{d_w}$ extracts features from the input data, and $f_w : \mathbb{R}^{d_w} \rightarrow \mathcal{Y}$ is fine-tuned using supervised data with ground truth labels. The strong model, on the other hand, aims to learn a mapping $f_{sw} \circ h_s$, where $h_s : \mathcal{X} \rightarrow \mathbb{R}^{d_s}$ is the representation, and $f_{sw} \in \mathcal{F}_s$ is a task-specific function from a hypothesis class $\mathcal{F}_s : \mathbb{R}^{d_s} \rightarrow \mathcal{Y}$. The strong model leverages the representation h_s to improve performance on the fine-tuning task. In the convention setting of AI alignment (Ouyang et al., 2022), the model is learned through human-annotated ground truth data:

$$f_s = \operatorname{argmin}_{f \in \mathcal{F}_s} d_{\mathcal{P}}(f^* \circ h^*, f \circ h_s). \quad (5)$$

Nevertheless, the acquisition of human-generated data is both costly and time-consuming. To address this challenge, the WTSG framework leverages weak supervision from the weak model's predictions, enabling the strong model to be trained through population risk minimization:

$$f_{sw} = \operatorname{argmin}_{f \in \mathcal{F}_s} d_{\mathcal{P}}(f_w \circ h_w, f \circ h_s). \quad (6)$$

In practice, we label n i.i.d. samples using the weak model and minimize the empirical risk to conduct WTSG:

$$\hat{f}_{sw} = \operatorname{argmin}_{f \in \mathcal{F}_s} \hat{d}_{\mathcal{P}}(f_w \circ h_w, f \circ h_s). \quad (7)$$

Denote the labeling function $F^* = f^* \circ h^*$, strong ceiling model $F_s = f_s \circ h_s$, weak model $F_w = f_w \circ h_w$, and strong models $F_{sw} = f_{sw} \circ h_s$, $\hat{F}_{sw} = \hat{f}_{sw} \circ h_s$, respectively.

4 UNIVERSAL RESULTS IN WTSG

In this section, we consider the classification problem, where $d_{\mathcal{P}}$ is the KL divergence loss defined in Equation (2). We first establish lower and upper generalization error bounds of the strong model in WTSG in Section 4.1. Then the lower and upper calibration error bounds are shown in Section 4.2.

4.1 LOWER AND UPPER BOUND

Theorem 1 (Proved in Appendix B.1). *Given the data domain $\mathcal{X}$, output domain $\mathcal{Y}$ and models F_{sw}, F_w, F^* defined above. Then there holds*

$$d_{\mathcal{P}}(F^*, F_w) - C_1 \sqrt{d_{\mathcal{P}}(F_w, F_{sw})} \leq d_{\mathcal{P}}(F^*, F_{sw}) \leq d_{\mathcal{P}}(F^*, F_w) + C_1 \sqrt{d_{\mathcal{P}}(F_w, F_{sw})}, \quad (8)$$

where C_1 is a positive constant.

Remark. *The proof can be also extended to the regression setting, which is provided in Appendix B.2.*

Theorem 1 provides a quantitative framework for assessing the performance gap between weak model and strong model in WTSG. Specifically, the value of $d_{\mathcal{P}}(F^*, F_{sw})$ is constrained by two terms: (1) $d_{\mathcal{P}}(F^*, F_{sw})$, which reflects the performance of the weak model, and (2) $d_{\mathcal{P}}(F_w, F_{sw})$, which is decided by the optimization result in Equation (6) and measures how the strong model learns to imitate the weak supervisor. This result is examined from two complementary perspectives: a lower bound and an upper bound. They offer insights into the fundamental limitation and theoretical guarantee for WTSG.

Lower bound. The lower bound indicates the fundamental limitation: $d_{\mathcal{P}}(F^*, F_{sw})$ cannot be arbitrarily reduced. Firstly, a minimal $d_{\mathcal{P}}(F^*, F_{sw})$ is intrinsically tied to the weak model performance $d_{\mathcal{P}}(F^*, F_w)$. To improve the strong model, the weak model becomes critical—that is, $d_{\mathcal{P}}(F^*, F_w)$ should be as small as possible. It underscores the importance of **carefully selecting the weak model** (Burns et al., 2023; Yang et al., 2025). Secondly, the performance improvement of strong model over the weak model cannot exceed $\mathcal{O}\left(\sqrt{d_{\mathcal{P}}(F_w, F_{sw})}\right)$. In WTSG, while the student-supervisor disagreement $d_{\mathcal{P}}(F_w, F_{sw})$ is minimized in Equation (6), we anticipate $\mathcal{O}\left(\sqrt{d_{\mathcal{P}}(F_w, F_{sw})}\right)$ to remain relatively small. However, a paradox arises: achieving a smaller $d_{\mathcal{P}}(F^*, F_{sw})$ necessitates a larger $d_{\mathcal{P}}(F_w, F_{sw})$. This implies that **the performance improvement of WTSG is probably constrained by its own optimization objective**.

Upper bound. The upper bound provides a theoretical guarantee for WTSG by ensuring that $d_{\mathcal{P}}(F^*, F_{sw})$ remains bounded and does not grow arbitrarily large. Firstly, effective WTSG requires choosing a weak model that produces supervision signal closely aligned with the true score, i.e., achieving a small $d_{\mathcal{P}}(F^*, F_w)$. To this end, employing a stronger weak model is crucial to obtain a tighter upper bound of $d_{\mathcal{P}}(F^*, F_{sw})$. Secondly, the worst-case performance of the strong model is constrained by the sum of $d_{\mathcal{P}}(F^*, F_w)$ and $\mathcal{O}\left(\sqrt{d_{\mathcal{P}}(F_w, F_{sw})}\right)$. By appropriately selecting the weak model and determining the minimizer of Equation (6), both $d_{\mathcal{P}}(F^*, F_w)$ and $\mathcal{O}\left(\sqrt{d_{\mathcal{P}}(F_w, F_{sw})}\right)$ can be kept small, ensuring the practicality of the strong model.

4.2 CALIBRATION IN WEAK-TO-STRONG GENERALIZATION

In this section, we further explore WTSG through the lens of calibration (Kumar et al., 2019), which requires that the prediction confidence should match the actual outcome. We first state the definition of Marginal Calibration Error (MCE) (Kumar et al., 2019), which is an extended version of Expected Calibration Error (ECE) (Guo et al., 2017) designed for multi-class classification. In particular, we use an ℓ_1 version of it, with the weight constant $\frac{1}{k}$ omitted.

Definition 1 (Marginal Calibration Error (Kumar et al., 2019)). Let $x \in \mathcal{X}$, ground truth $y = [y_1, \dots, y_k]^T \in \{0, 1\}^k$ where $\sum_{i=1}^k y_i = 1$, and a model $f : \mathcal{X} \rightarrow \mathcal{Y}$. Define the marginal calibration error of f as:

$$MCE(f) = \sum_{i=1}^k \mathbb{E}_x |[f(x)]_i - \mathbb{P}[y_i = 1 | [f(x)]_i]|. \quad (9)$$

It measures the difference between model confidence and actual outcome, and $MCE(f) \in [0, 2]$. For binary classification, MCE is twice the ECE . We shed light on upper and lower bounds of calibration of the strong model.

Theorem 2 (Proved in Appendix B.3). Let $MCE(\cdot)$ be the marginal calibration error in Definition 1. Then there holds

$$|MCE(F_{sw}) - MCE(F_w)| \leq 2 \cdot \sqrt{1 - \exp(-d_{\mathcal{P}}(F_w, F_{sw}))}. \quad (10)$$

Theorem 2 demonstrates that the calibration error of F_{sw} is influenced by two key factors: (1) the calibration error of F_w , and (2) the teacher-student disagreement, as characterized by the optimization result in Equation (6). This theoretical result yields two insights. First, to achieve a strong model with acceptable calibration, the weak teacher should also exhibit acceptable calibration. Otherwise, the strong model will inherit a non-trivial calibration error from the weak teacher as $d_{\mathcal{P}}(F_w, F_{sw})$ goes to zero. Second, closely imitating the weak supervisor minimizes $d_{\mathcal{P}}(F_w, F_{sw})$, causing the calibration errors of the strong and weak models to converge. Taking them together, to ensure WTSG with reasonable calibration and prevent a poorly-calibrated F_{sw} , it is crucial to **avoid using a poorly-calibrated weak model with an overfitted strong model**. Additionally, since models with larger capacity may exhibit higher calibration errors (Guo et al., 2017), **a potential trade-off exists between the weak model’s calibration error and the teacher-student disagreement**. In other words, $MCE(F_w)$ and $\sqrt{1 - \exp(-d_{\mathcal{P}}(F_w, F_{sw}))}$ may not be minimized simultaneously, posing a challenge in selecting the weak model and designing an effective optimization strategy to achieve better calibration in the strong model.

4.3 EXPERIMENTAL VALIDATION IN LANGUAGE MODELS

In this section, we use language models to verify our theoretical results in WTSG.

4.3.1 EXPERIMENTAL SETTING

Dataset. We define the alignment objective as enabling a weak model to guide a strong model in achieving harmlessness. To this end, we employ CAI-Harmless (Bai et al., 2022b), which is a widely adopted single-turn harmless dataset for reward modeling task. Each sample is structured as $(x; y_c, y_r)$, where x denotes the prompt, and y_c and y_r represent the human-preferred and human-rejected completions, respectively. The dataset is randomly split into three 4K-sample subsets: one for fine-tuning both weak and strong base models, another for weak supervision via weak model predictions, and the last for testing and evaluation.

Model. To explore weak-to-strong generalization, we utilize GPT-2 series (Radford et al., 2019) (including GPT-2-Base, GPT-2-Medium, GPT-2-Large, and GPT-2-XL) and Pythia series (Biderman et al., 2023) (including Pythia-70M, Pythia-160M, Pythia-410M and Pythia-1B). For each model, we append a linear projection head to facilitate logit predictions for each completion pair $\tilde{x} = (x; y_c, y_r)$. Consequently, the task can be framed as a binary classification problem, where the model F predicts the soft label as

$$F(\tilde{x}) = \text{Sigmoid}(F(y_c) - F(y_r)).$$

Training. The models are trained via KL divergence loss. Training details are in Appendix C.1.

Metric. To evaluate whether a model F can effectively distinguish between chosen and rejected completions (y_c and y_r) for a given prompt x , we aim for F to assign a higher score to the chosen completion compared to the rejected one. Specifically, this requires $F(y_c) - F(y_r) > 0$ for each completion pair $\tilde{x} = (x; y_c, y_r)$, which implies $F(\tilde{x}) > 0.5$. Accordingly, the test accuracy of a model F is reported as the fraction of predictions that satisfy $F(\tilde{x}) > 0.5$.

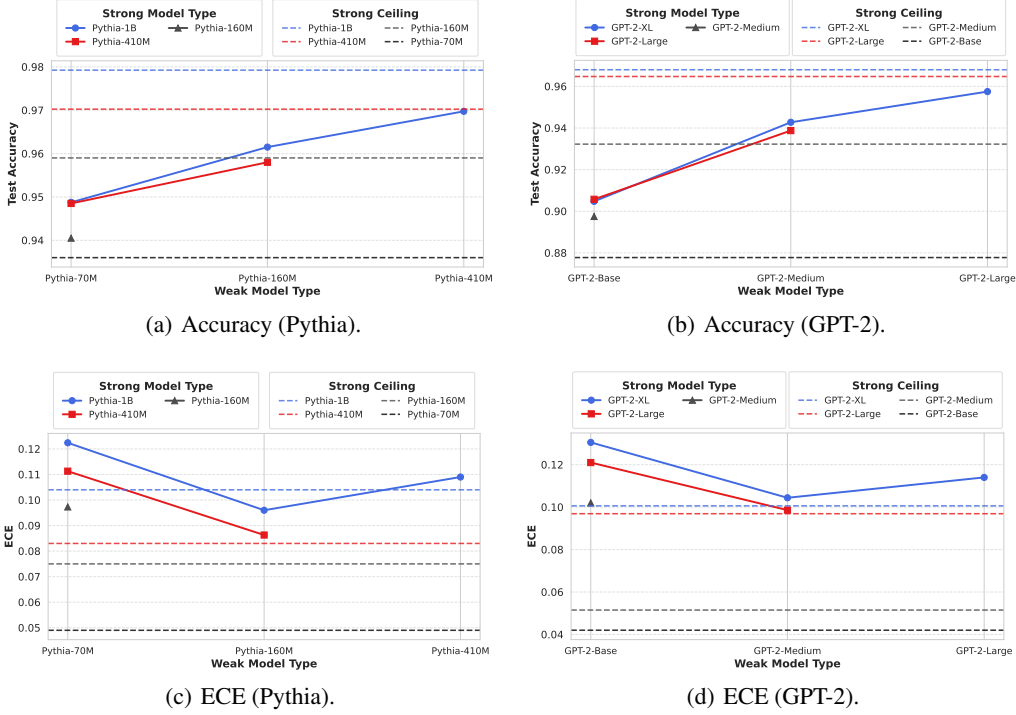


Figure 1: Accuracy and calibration results for Pythia and GPT-2 series. **(a)** Test accuracies of Pythia series. Each curve demonstrates the variation in accuracy of WTSG as strong models are supervised by weak models of varying capabilities. “Strong Ceiling” corresponds to models fine-tuned using ground truth data. **(b)** Test accuracies of GPT-2 series. **(c)** Expected calibration errors of Pythia series. Each curve depicts the change in ECE as strong models are supervised by different weak teachers. **(d)** Expected calibration errors of GPT-2 series.

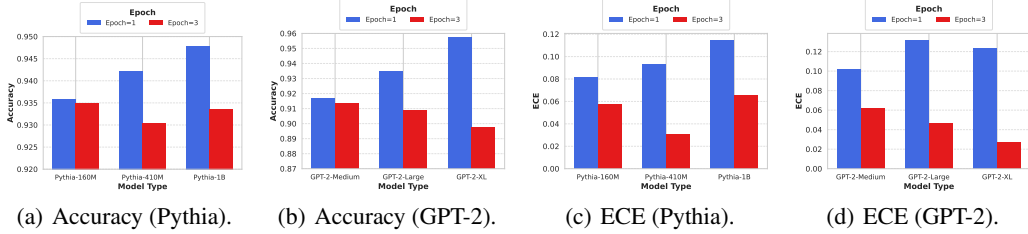


Figure 2: Ablation study for the Pythia and GPT-2 series. **(a)-(b)** Test accuracies of Pythia and GPT-2. The accuracies of Pythia-70M and GPT-2-Base fine-tuned on ground truth data is 92.45% and 90.95%, respectively. **(c)-(d)** ECE of Pythia and GPT-2. The ECE of Pythia-70M and GPT-2-Base fine-tuned on ground truth data is 0.049 and 0.042, respectively.

4.3.2 RESULTS AND ANALYSIS

The main results of WTSG for Pythia and GPT-2 series are shown in Figure 1. To further investigate the optimization result $d_{\mathcal{P}}(F_w, F_{sw})$ in WTSG, we increase the number of epochs to train a strong model that more closely imitates the weak model. The corresponding results are in Figure 2.

Main results. Figure 1(a) and Figure 1(b) demonstrate that, for the same strong model, the generalization of WTSG increases when supervised by a weak model of greater capacity. It is consistent with Theorem 1. Figure 1(c) and Figure 1(d) illustrate the results for calibration errors. Interestingly, we observe that stronger models with higher capacity tend to exhibit larger ECE. Furthermore, increasing the weak model’s capacity results in a U-shaped trend in ECE. This pattern suggests a potential trade-off between the weak model’s calibration quality and the teacher-student disagreement, which is consistent with Theorem 2.

Ablation study. To investigate WTSG over extended training epochs, we design a series of teacher-student pairs with increasing model capacities. Specifically, we employ Pythia-70M as the weak teacher to supervise larger student models, including Pythia-160M, 410M, and 1B. We also utilize GPT-2-Base as the weak teacher for supervising GPT-2-Medium, Large, and XL. Figure 2 illustrates that as we increase the number of epochs to reduce $d_{\mathcal{P}}(F_w, F_{sw})$, *there is a simultaneous decline in both the accuracy and calibration error of other strong models*. Taking the Pythia series as an example, Figure 1(a) and Figure 1(c) demonstrate that Pythia-70M achieves the lowest accuracy and best ECE performance among the Pythia models. While Theorem 2 indicates that reducing $d_{\mathcal{P}}(F_w, F_{sw})$ causes the accuracy and calibration results of strong models to converge toward those of the weak model, our experiments show that increasing the number of epochs leads to reduced accuracy and ECE for Pythia-160M, 410M, and 1B. In other words, the accuracy and ECE of strong models approach those of the weak model, consistent with Theorem 1 and Theorem 2. And this trend is also observed in the GPT-2 series.

Potential overfitting. As the number of epochs increases, *the accuracy of GPT-2-XL drops even below that of GPT-2-Base (90.95%)*. This is attributed to the strong expressive power of GPT-2-XL, which leads to overfitting to the weak supervision provided by GPT-2-Base. Note that the upper bounds derived in Theorem 1 and Theorem 2 do not guarantee that the strong model will outperform the weak model in terms of both generalization performance and calibration properties. The underlying intuition is that if a strong model overfits to the weak supervision, it may closely mimic the weak model’s generalization and calibration behavior. Consequently, the strong model could end up performing on par with or potentially even worse than the weak model.

5 RESULTS BEYOND SQUARED LOSS

In regression problems under some assumptions, Charikar et al. (2024) proves that the strong model’s error is smaller than the weak model’s, with the gap at least the strong model’s error on the weak labels. This observation naturally raises the following question: *Can their proof be extended from squared loss to output distribution divergence?* In this section, we show how to theoretically bridge the gap between squared loss and KL divergence within the overall proof framework established in Charikar et al. (2024). To begin with, we restate an assumption used in previous study.

Assumption 1 (Convexity Assumption (Charikar et al., 2024)). *The strong model learns fine-tuning tasks from a function class $\mathcal{F}_s$, which is a convex set.*

It requires that, for any $f, g \in \mathcal{F}_s$, and for any $\lambda \in [0, 1]$, there exists $h \in \mathcal{F}_s$ such that for all $z \in \mathbb{R}^{d_s}$, $h(z) = \lambda f(z) + (1 - \lambda)g(z)$. To satisfy the convex set assumption, $\mathcal{F}_s$ can be the class of all linear functions. In these cases, $\mathcal{F}_s$ is a convex set. Note that it is validated by practice: a popular way to fine-tune a pre-trained model on task-specific data is by tuning the weights of only the last linear layer of the model (Howard & Ruder, 2018; Kumar et al., 2022).

5.1 UPPER BOUND (REALIZABILITY)

Firstly, we consider the case where $\exists f_s \in \mathcal{F}_s$ such that $F_s = F^*$ (also called “Realizability” (Charikar et al., 2024)). It means we can find a f_s such that $f_s \circ h_s = f^* \circ h^*$. This assumption implicitly indicates the strong power of pre-training. It requires that the representation h_s has learned extremely enough information during pre-training, which is reasonable in modern large language models pre-trained on very large corpus (Touvron et al., 2023; Achiam et al., 2023). The scale and diversity of the corpus ensure that the model is exposed to a broad spectrum of lexical, syntactic, and semantic structures, enhancing its ability to generalize effectively across varied language tasks.

We state our result in the realizable setting, which corresponds to Theorem 1 in Charikar et al. (2024).

Theorem 3 (Proved in Appendix B.4). *Given F^* , F_w and F_{sw} defined above. Consider $\mathcal{F}_s$ that satisfies Assumption 1. Consider WTSG using reverse KL divergence loss:*

$$f_{sw} = \operatorname{argmin}_{f \in \mathcal{F}_s} d_{\mathcal{P}}(f \circ h_s, f_w \circ h_w).$$

Assume that $\exists f_s \in \mathcal{F}_s$ such that $F_s = F^$. Then*

$$d_{\mathcal{P}}(F^*, F_{sw}) \leq d_{\mathcal{P}}(F^*, F_w) - d_{\mathcal{P}}(F_{sw}, F_w). \quad (11)$$

Remark. *The corresponding theorem and proof in the case of forward KL divergence loss is provided in Corollary 1 from Appendix B.5, under an additional assumption.*

In contrast to the symmetric squared loss studied in prior work (Charikar et al., 2024), the emergence of the reverse KL divergence is inherently tied to the asymmetric properties of the KL divergence. Although extending previous work to both forward and reverse KL divergences presents significant technical challenges, our results demonstrate the theoretical guarantees of WTSG in these settings. In Inequality (11), the left-hand side represents the error of the weakly-supervised strong model on the true data. On the right-hand side, the first term denotes the true error of the weak model, while the second term captures the disagreement between the strong and weak models, which also serves as the minimization objective in WTSG. This inequality indicates that the weakly-supervised strong model improves upon the weak model by at least the magnitude of their disagreement, $d_{\mathcal{P}}(F_{sw}, F_w)$. To reduce the error of F_{sw} , Theorem 3 aligns with Theorem 1, highlighting the importance of selecting an effective weak model and the inherent limitations of the optimization objective in WTSG.

5.2 UPPER BOUND (NON-REALIZABILITY)

Now we relax the “realizability” condition and draw n i.i.d. samples to perform WTSG. We provide the result in the “unrealizable” setting, where the condition $F_s = F^*$ may not be satisfied for any $f_s \in \mathcal{F}_s$. It corresponds to Theorem 2 in Charikar et al. (2024).

Theorem 4 (Proved in Appendix B.6). *Given F^* , F_w and F_{sw} defined above. Consider $\mathcal{F}_s$ that satisfies Assumption 1. Consider weak-to-strong generalization using reverse KL:*

$$\begin{aligned} f_{sw} &= \operatorname{argmin}_{f \in \mathcal{F}_s} d_{\mathcal{P}}(f \circ h_s, f_w \circ h_w), \\ \hat{f}_{sw} &= \operatorname{argmin}_{f \in \mathcal{F}_s} \hat{d}_{\mathcal{P}}(f \circ h_s, f_w \circ h_w), \end{aligned}$$

Denote $d_{\mathcal{P}}(F^*, F_s) = \varepsilon$. With probability at least $1 - \delta$ over the draw of n i.i.d. samples, there holds

$$d_{\mathcal{P}}(F^*, \hat{F}_{sw}) \leq d_{\mathcal{P}}(F^*, F_w) - d_{\mathcal{P}}(\hat{F}_{sw}, F_w) + \mathcal{O}(\sqrt{\varepsilon}) + \mathcal{O}\left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}_s}}{n}}\right) + \mathcal{O}\left(\sqrt{\frac{\log(1/\delta)}{n}}\right), \quad (12)$$

where $\mathcal{C}_{\mathcal{F}_s}$ is a constant capturing the complexity of the function class $\mathcal{F}_s$, and the asymptotic notation is with respect to $\varepsilon \rightarrow 0, n \rightarrow \infty$.

Remark. The extension to forward KL divergence loss is provided in Corollary 2 from Appendix B.7, under an additional assumption.

Compared to Inequality (11), this bound introduces two another error terms: the first term of $\mathcal{O}(\sqrt{\varepsilon})$ arises due to the non-realizability assumption, and diminishes as the strong ceiling model F_s becomes more expressive. The remaining two error terms arise from the strong model $\hat{F}_{sw}$ being trained on a finite weakly-labeled sample. They also asymptotically approach zero as the sample size increases.

5.3 SYNTHETIC EXPERIMENTS

In this section, we conduct experiments on synthetically generated data to validate the theoretical results in Section 5. While drawing inspiration from the theoretical framework of Charikar et al. (2024), we extend their synthetic experiments by replacing the squared loss used in their work with the output distribution divergence defined in Equation (4).

5.3.1 EXPERIMENTAL SETTING

In our setup, The data distribution $\mathcal{P}$ is chosen as $\mathcal{N}(0, \sigma^2 I)$, with $\sigma = 500$ to ensure the data is well-dispersed. The ground truth representation $h^* : \mathbb{R}^8 \rightarrow \mathbb{R}^{16}$ is implemented as a randomly initialized 16-layer multi-layer perceptron (MLP) with ReLU activations. Let the weak model and strong model representations $h_w, h_s : \mathbb{R}^8 \rightarrow \mathbb{R}^{16}$ be 2-layer and 8-layer MLP with ReLU activations, respectively. Given h_w and h_s frozen, both the strong and weak models learn from the fine-tuning task class $\mathcal{F}_s$, which consists of linear functions mapping $\mathbb{R}^{16} \rightarrow \mathbb{R}$. This makes $\mathcal{F}_s$ a convex set.

For the “realizable” setting, we set $h_s = h^*$. For the “unrealizable” setting, we adopt the approach of Charikar et al. (2024) and investigate two methods for generating weak and strong representations: (1) **Pre-training**: 20 models $f_1^*, \dots, f_{20}^* : \mathbb{R}^8 \rightarrow \mathbb{R}^{16}$ are randomly sampled as fine-tuning tasks. 2000 data points are independently generated from $\mathcal{P}$ for these tasks. Accordingly, h_w and h_s

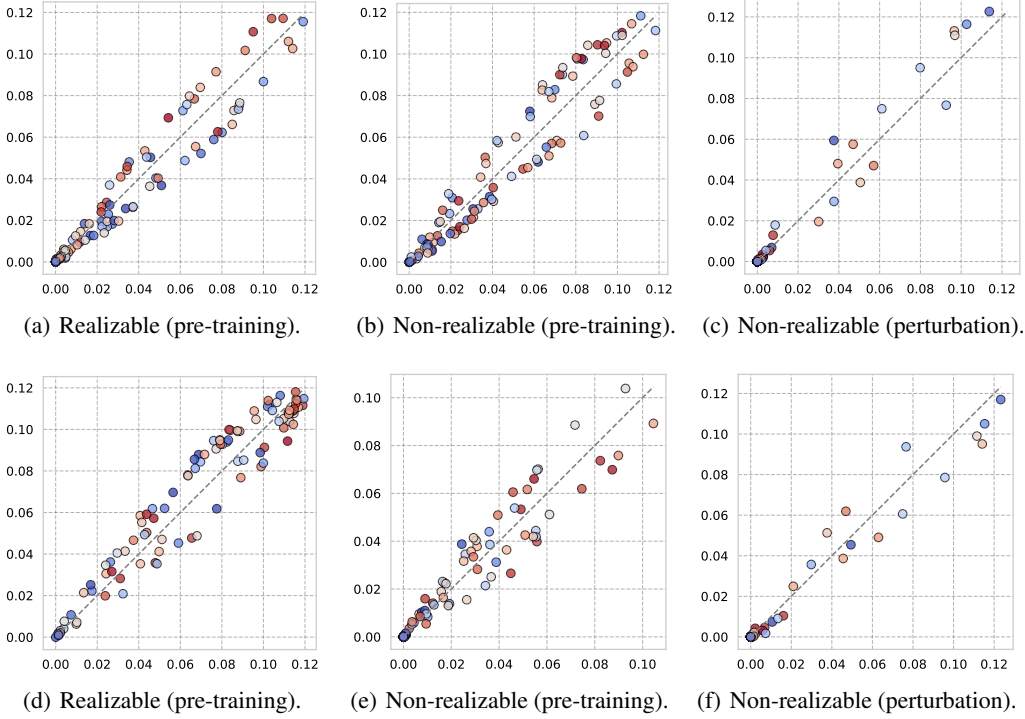


Figure 3: Experiments on synthetic data using reverse KL divergence loss (a-c) and forward KL divergence loss (d-f). Each point corresponds to a task and the gray dotted line represents $y = x$. h^* is a 16-layer MLP. (a,d) Realizable (pre-training): $h_s = h^*$, and h_w is a 2-layer MLP obtained by pre-training. (b,e) Non-realizable (pre-training): h_s is an 8-layer MLP, and h_w is a 2-layer MLP. Both h_s and h_w are obtained by pre-training. (c,f) Non-realizable (perturbation): Both h_s and h_w are obtained by directly perturbing the weights in h^* : $h_s = h^* + \mathcal{N}(0, 0.01)$, and $h_w = h^* + \mathcal{N}(0, 9)$.

are obtained by minimizing the average output distribution divergence between ground truth label ($f_t^* \circ h^*$) and model prediction ($f_t^* \circ h_w$ and $f_t^* \circ h_s$) over the 20 tasks. (2) **Perturbations:** As an alternative, we directly perturb the parameters of h^* to obtain the weak and strong representations. Specifically, we add independent Gaussian noises $\mathcal{N}(0, \sigma_s^2)$ and $\mathcal{N}(0, \sigma_w^2)$ to every parameter in h^* to generate h_s and h_w , respectively. To ensure the strong representation h_s is closer to h^* than h_w (Charikar et al., 2024), we set $\sigma_s = 0.1$ and $\sigma_w = 3$.

Weak Model Finetuning. We freeze the weak model representation h_w and train the weak models on new fine-tuning tasks. We randomly sample 100 new fine-tuning tasks $f_{21}^*, \dots, f_{120}^* : \mathbb{R}^8 \rightarrow \mathbb{R}^{16}$, and independently generate another 2000 data points from $\mathcal{P}$. For each task $t \in \{21, \dots, 120\}$, the corresponding weak model is obtained by minimizing the output distribution divergence between ground truth label and weak model prediction.

Weak-to-Strong Supervision. Using the trained weak models, we generate weakly labeled data to supervise the strong model. Specifically, we first independently generate another 2000 data points from $\mathcal{P}$. Then for each task $t \in \{21, \dots, 120\}$, the strong model is obtained by minimizing the output distribution divergence between weak model supervision and strong model prediction. At this stage, the weak-to-strong training procedure is complete. The detailed introduction of above is in Appendix C.2.

Evaluation. We independently draw an additional 2000 samples from $\mathcal{P}$ to construct the test set. They are used to estimate $d_{\mathcal{P}}(F^*, F_{sw})$, $d_{\mathcal{P}}(F^*, F_w)$ and $d_{\mathcal{P}}(F_{sw}, F_w)$ for each task $t \in \{21, \dots, 120\}$. We estimate these quantities using their empirical counterparts: $\hat{d}_{\mathcal{P}}(F^*, F_w)$, $\hat{d}_{\mathcal{P}}(F^*, F_{sw})$, and $\hat{d}_{\mathcal{P}}(F_{sw}, F_w)$. To validate Theorem 3-4 and visualize the trend clearly, we plot $\hat{d}_{\mathcal{P}}(F^*, F_w) - \hat{d}_{\mathcal{P}}(F^*, F_{sw})$ on the x -axis versus $\hat{d}_{\mathcal{P}}(F_{sw}, F_w)$ on the y -axis. The results are presented in Figure 3(a)-(c). We also examine forward KL divergence loss. To validate Corollary 1-2, which extend Theorem 3-

4 to the case of using forward KL divergence loss in WTSG, we plot $\hat{d}_{\mathcal{P}}(F_w, F^*) - \hat{d}_{\mathcal{P}}(F_{sw}, F^*)$ on the x -axis versus $\hat{d}_{\mathcal{P}}(F_w, F_{sw})$ on the y -axis. The results are presented in Figure 3(d)-(f).

5.3.2 RESULTS AND ANALYSIS

Reverse KL divergence loss. Similar to previous results (Charikar et al., 2024), the points in our experiments also cluster around the line $y = x$. This suggests that $\hat{d}_{\mathcal{P}}(F^*, F_w) - \hat{d}_{\mathcal{P}}(F^*, F_{sw}) \approx \hat{d}_{\mathcal{P}}(F_{sw}, F_w)$. It is consistent with our theoretical framework, suggesting that the improvement over the weak teacher can be quantified by the disagreement between strong and weak models.

Forward KL divergence loss. The observed trend closely mirrors that of reverse KL. The dots are generally around the line $y = x$. It suggest that the relationship $\hat{d}_{\mathcal{P}}(F_w, F^*) - \hat{d}_{\mathcal{P}}(F_{sw}, F^*) \approx \hat{d}_{\mathcal{P}}(F_w, F_{sw})$ may also hold, indicating a similar theoretical guarantee for forward KL in WTSG.

6 CONCLUSION

This paper provides a comprehensive theoretical framework for understanding the capability and limitation of weak-to-strong generalization. In the classification setting, we establish upper and lower bounds for both the generalization and calibration errors of the strong model, revealing that the primary limitations arise from the weak model and the optimization objective. These bounds emphasize two critical insights: (1) the weak model must demonstrate strong generalization and calibration performance, and (2) the strong model should avoid excessive training to prevent overfitting on weak supervision. In the regression setting, we extended prior work to output distribution divergence loss, proving that a strong model can outperform its weak teacher by at least their disagreement under certain assumptions. Our theoretical findings were validated through experiments with language models and MLP, providing practical insights and some interesting observations for the WTSG paradigm. Overall, we hope this work enhances the understanding of weak-to-strong generalization and encourages future research to unlock its promise for human-aligned AI systems.

BROADER IMPACT AND ETHICS STATEMENT

This work on weak-to-strong generalization aims to improve the alignment of superhuman models with human values. While our theoretical and empirical insights highlight the potential of this approach, we acknowledge the risks of propagating biases or errors from the weak model to the strong model. To address these concerns, we emphasize the importance of ensuring the weak model’s generalization and calibration, as well as carefully balancing the strong model’s optimization to avoid over-reliance on weak supervision. We encourage rigorous testing, transparency, and ongoing monitoring in real-world applications to ensure the safe and ethical deployment of such systems. Our work contributes to the broader effort of aligning advanced AI with human values, but its implementation must prioritize fairness, accountability, and safety.

ACKNOWLEDGEMENTS

We are deeply grateful to Shaojie Li, Chenyu Zheng, Gengze Xu for their constructive suggestions. This research was supported by National Natural Science Foundation of China (No.62476277), National Key Research and Development Program of China(NO. 2024YFE0203200), CCF-ALIMAMA TECH Kangaroo Fund(No.CCF-ALIMAMA OF 2024008), and Huawei-Renmin University joint program on Information Retrieval. We also acknowledge the support provided by the fund for building worldclass universities (disciplines) of Renmin University of China and by the funds from Beijing Key Laboratory of Big Data Management and Analysis Methods, Gaoling School of Artificial Intelligence, Renmin University of China, from Engineering Research Center of Next-Generation Intelligent Search and Recommendation, Ministry of Education, from Intelligent Social Governance Interdisciplinary Platform, Major Innovation & Planning Interdisciplinary Platform for the “DoubleFirst Class” Initiative, Renmin University of China, from Public Policy and Decision-making Research Lab of Renmin University of China, and from Public Computing Cloud, Renmin University of China.

REFERENCES

- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- Aakriti Agrawal, Mucong Ding, Zora Che, Chenghao Deng, Anirudh Satheesh, John Langford, and Furong Huang. Ensemw2s: Can an ensemble of llms be leveraged to obtain a stronger llm? *arXiv preprint arXiv:2410.04571*, 2024.
- Gholamali Aminian, Mahed Abroshan, Mohammad Mahdi Khalili, Laura Toni, and Miguel Rodrigues. An information-theoretical approach to semi-supervised learning under covariate-shift. In *International Conference on Artificial Intelligence and Statistics*, pp. 7433–7449, 2022.
- Yuntao Bai, Andy Jones, Kamal Ndousse, Amanda Askell, Anna Chen, Nova DasSarma, Dawn Drain, Stanislav Fort, Deep Ganguli, Tom Henighan, et al. Training a helpful and harmless assistant with reinforcement learning from human feedback. *arXiv preprint arXiv:2204.05862*, 2022a.
- Yuntao Bai, Saurav Kadavath, Sandipan Kundu, Amanda Askell, Jackson Kernion, Andy Jones, Anna Chen, Anna Goldie, Azalia Mirhoseini, Cameron McKinnon, et al. Constitutional ai: Harmlessness from ai feedback. *arXiv preprint arXiv:2212.08073*, 2022b.
- Peter L Bartlett and Shahar Mendelson. Rademacher and gaussian complexities: Risk bounds and structural results. *Journal of Machine Learning Research*, 3:463–482, 2002.
- Lucas Beyer, Xiaohua Zhai, Amélie Royer, Larisa Markeeva, Rohan Anil, and Alexander Kolesnikov. Knowledge distillation: A good teacher is patient and consistent. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pp. 10925–10934, 2022.
- Stella Biderman, Hailey Schoelkopf, Quentin Gregory Anthony, Herbie Bradley, Kyle O’Brien, Eric Hallahan, Mohammad Aflah Khan, Shivanshu Purohit, USVSN Sai Prashanth, Edward Raff, et al. Pythia: A suite for analyzing large language models across training and scaling. In *International Conference on Machine Learning*, pp. 2397–2430, 2023.
- Yuheng Bu, Shaofeng Zou, and Venugopal V Veeravalli. Tightening mutual information-based bounds on generalization error. *IEEE Journal on Selected Areas in Information Theory*, 1(1): 121–130, 2020.
- Collin Burns, Pavel Izmailov, Jan Hendrik Kirchner, Bowen Baker, Leo Gao, Leopold Aschenbrenner, Yining Chen, Adrien Ecoffet, Manas Joglekar, Jan Leike, et al. Weak-to-strong generalization: Eliciting strong capabilities with weak supervision. *arXiv preprint arXiv:2312.09390*, 2023.
- Stephen Casper, Xander Davies, Claudia Shi, Thomas Krendl Gilbert, Jérémy Scheurer, Javier Rando, Rachel Freedman, Tomasz Korbak, David Lindner, Pedro Freire, et al. Open problems and fundamental limitations of reinforcement learning from human feedback. *Transactions on Machine Learning Research*, 2023.
- Moses Charikar, Chirag Pabbaraju, and Kirankumar Shiragur. Quantifying the gain in weak-to-strong generalization. *Advances in neural information processing systems*, 2024.
- Qi Chen, Changjian Shui, and Mario Marchand. Generalization bounds for meta-learning: An information-theoretic analysis. *Advances in Neural Information Processing Systems*, 34:25878–25890, 2021.
- Yangyi Chen, Lifan Yuan, Ganqu Cui, Zhiyuan Liu, and Heng Ji. A close look into the calibration of pre-trained language models. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 1343–1367, 2023.
- Lele Cheng, Xiangzeng Zhou, Liming Zhao, Dangwei Li, Hong Shang, Yun Zheng, Pan Pan, and Yinghui Xu. Weakly supervised learning with side information for noisy labeled images. In *The European Conference on Computer Vision*, pp. 306–321, 2020.
- P Chu and D Messerschmitt. A frequency weighted itakura-saito spectral distance measure. *IEEE Transactions on Acoustics, Speech, and Signal Processing*, 30(4):545–560, 1982.

- Thomas M Cover. *Elements of information theory*. John Wiley & Sons, 1999.
- Shrey Desai and Greg Durrett. Calibration of pre-trained transformers. *arXiv preprint arXiv:2003.07892*, 2020.
- Inderjit S. Dhillon. Learning with bregman divergences. <https://www.cs.utexas.edu/~inderjit/Talks/bregtut.pdf>, 2007.
- Monroe D Donsker and SR Srinivasa Varadhan. Asymptotic evaluation of certain markov process expectations for large time. iv. *Communications on pure and applied mathematics*, 36(2):183–212, 1983.
- Cédric Févotte, Nancy Bertin, and Jean-Louis Durrieu. Nonnegative matrix factorization with the itakura-saito divergence: With application to music analysis. *Neural computation*, 21(3):793–830, 2009.
- Jiahui Geng, Fengyu Cai, Yuxia Wang, Heinz Koepl, Preslav Nakov, and Iryna Gurevych. A survey of confidence estimation and calibration in large language models. In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pp. 6577–6595, 2024.
- Chuan Guo, Geoff Pleiss, Yu Sun, and Kilian Q Weinberger. On calibration of modern neural networks. In *International Conference on Machine Learning*, pp. 1321–1330, 2017.
- Han Guo, Ramakanth Pasunuru, and Mohit Bansal. An overview of uncertainty calibration for text classification and the role of distillation. In *Proceedings of the 6th Workshop on Representation Learning for NLP (RepL4NLP-2021)*, pp. 289–306, 2021.
- Jianyuan Guo, Hanling Chen, Chengcheng Wang, Kai Han, Chang Xu, and Yunhe Wang. Vision superalignment: Weak-to-strong generalization for vision foundation models. *arXiv preprint arXiv:2402.03749*, 2024.
- Yue Guo and Yi Yang. Improving weak-to-strong generalization with reliability-aware alignment. *arXiv preprint arXiv:2406.19032*, 2024.
- Fredrik Hellström, Giuseppe Durisi, Benjamin Guedj, and Maxim Raginsky. Generalization bounds: Perspectives from information theory and pac-bayes. *arXiv preprint arXiv:2309.04381*, 2023.
- Geoffrey Hinton. Distilling the knowledge in a neural network. *arXiv preprint arXiv:1503.02531*, 2015.
- Jeremy Howard and Sebastian Ruder. Universal language model fine-tuning for text classification. In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 328–339, 2018.
- Fumitada Itakura. Analysis synthesis telephony based on the maximum likelihood method. *Reports of the 6-th Int. Cong. Acoust*, 1968.
- Jiaming Ji, Tianyi Qiu, Boyuan Chen, Borong Zhang, Hantao Lou, Kaile Wang, Yawen Duan, Zhonghao He, Jiayi Zhou, Zhaowei Zhang, et al. Ai alignment: A comprehensive survey. *arXiv preprint arXiv:2310.19852*, 2023.
- HyunJin Kim, Xiaoyuan Yi, Jing Yao, Jianxun Lian, Muhua Huang, Shitong Duan, JinYeong Bak, and Xing Xie. The road to artificial superintelligence: A comprehensive survey of superalignment. *arXiv preprint arXiv:2412.16468*, 2024.
- Volodymyr Kuleshov, Nathan Fenner, and Stefano Ermon. Accurate uncertainties for deep learning using calibrated regression. In *International conference on machine learning*, pp. 2796–2804, 2018.
- Meelis Kull, Miquel Perello Nieto, Markus Kängsepp, Telmo Silva Filho, Hao Song, and Peter Flach. Beyond temperature scaling: Obtaining well-calibrated multi-class probabilities with dirichlet calibration. *Advances in neural information processing systems*, 32, 2019.

- Ananya Kumar, Percy S Liang, and Tengyu Ma. Verified uncertainty calibration. *Advances in Neural Information Processing Systems*, 32, 2019.
- Ananya Kumar, Aditi Raghunathan, Robbie Jones, Tengyu Ma, and Percy Liang. Fine-tuning can distort pretrained features and underperform out-of-distribution. In *International Conference on Learning Representations*, 2022.
- Hunter Lang, David Sontag, and Aravindan Vijayaraghavan. Theoretical analysis of weak-to-strong generalization. *Advances in neural information processing systems*, 2024.
- Michel Ledoux and Michel Talagrand. *Probability in Banach Spaces: isoperimetry and processes*. Springer Science & Business Media, 2013.
- Sebastian Lee, Sebastian Goldt, and Andrew Saxe. Continual learning in the teacher-student setup: Impact of task similarity. In *International Conference on Machine Learning*, pp. 6109–6119, 2021.
- Shaojie Li, Bowei Zhu, and Yong Liu. Algorithmic stability unleashed: Generalization bounds with unbounded losses. In *Forty-first International Conference on Machine Learning*, 2024.
- Percy Liang, Rishi Bommasani, Tony Lee, Dimitris Tsipras, Dilara Soylu, Michihiro Yasunaga, Yian Zhang, Deepak Narayanan, Yuhuai Wu, Ananya Kumar, et al. Holistic evaluation of language models. *Transactions on Machine Learning Research*, 2023.
- Lydia T Liu, Max Simchowitz, and Moritz Hardt. The implicit fairness criterion of unconstrained learning. In *International Conference on Machine Learning*, pp. 4051–4060, 2019.
- Yuejiang Liu and Alexandre Alahi. Co-supervised learning: Improving weak-to-strong generalization with hierarchical mixture of experts. *arXiv preprint arXiv:2402.15505*, 2024.
- Tambet Matiisen, Avital Oliver, Taco Cohen, and John Schulman. Teacher–student curriculum learning. *IEEE transactions on neural networks and learning systems*, 31(9):3732–3740, 2019.
- Alireza Mehrtash, William M Wells, Clare M Tempny, Purang Abolmaesumi, and Tina Kapur. Confidence calibration and predictive uncertainty estimation for deep medical image segmentation. *IEEE transactions on medical imaging*, 39(12):3868–3878, 2020.
- Zhong Meng, Jinyu Li, Yong Zhao, and Yifan Gong. Conditional teacher-student learning. In *IEEE International Conference on Acoustics, Speech and Signal Processing*, pp. 6445–6449, 2019.
- Gabriel Meseguer-Brocal, Alice Cohen-Hadria, and Geoffroy Peeters. Dali: A large dataset of synchronized audio, lyrics and notes, automatically created using teacher-student machine learning paradigm. *arXiv preprint arXiv:1906.10606*, 2019.
- Mahdi Pakdaman Naeini, Gregory Cooper, and Milos Hauskrecht. Obtaining well calibrated probabilities using bayesian binning. In *Proceedings of the AAAI conference on artificial intelligence*, pp. 2901 – 2907, 2015.
- OpenAI. Introducing superalignment, 2024. URL <https://openai.com/index/introducing-superalignment/>.
- Maxime Oquab, Léon Bottou, Ivan Laptev, and Josef Sivic. Is object localization for free?-weakly-supervised learning with convolutional neural networks. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 685–694, 2015.
- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35:27730–27744, 2022.
- Dim P Papadopoulos, Jasper RR Uijlings, Frank Keller, and Vittorio Ferrari. Training object class detectors with click supervision. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 6374–6383, 2017.
- Martin Pawelczyk, Lillian Sun, Zhenting Qi, Aounon Kumar, and Himabindu Lakkaraju. Generalizing trust: Weak-to-strong trustworthiness in language models. *arXiv preprint arXiv:2501.00418*, 2024.

- Ankit Pensia, Varun Jog, and Po-Ling Loh. Generalization error bounds for noisy, iterative algorithms. In *2018 IEEE International Symposium on Information Theory*, pp. 546–550, 2018.
- Geoff Pleiss, Manish Raghavan, Felix Wu, Jon Kleinberg, and Kilian Q Weinberger. On fairness and calibration. *Advances in neural information processing systems*, 30, 2017.
- Dani Prasetyawan and Nakamoto Takamichi. Sensory evaluation of odor approximation using nmf with kullback-leibler divergence and itakura-saito divergence in mass spectrum space. *Journal of The Electrochemical Society*, 167(16):167520, 2020.
- Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, Ilya Sutskever, et al. Language models are unsupervised multitask learners. *OpenAI blog*, 1(8):9, 2019.
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model. *Advances in Neural Information Processing Systems*, 36, 2024.
- Alexander Ratner, Stephen H Bach, Henry Ehrenberg, Jason Fries, Sen Wu, and Christopher Ré. Snorkel: rapid training data creation with weak supervision. *The VLDB Journal*, 29(2):709–730, 2020.
- Rebecca Roelofs, Nicholas Cain, Jonathon Shlens, and Michael C Mozer. Mitigating bias in calibration error estimation. In *International Conference on Artificial Intelligence and Statistics*, pp. 4036–4054, 2022.
- Daniel Russo and James Zou. Controlling bias in adaptive data analysis using information theory. In *Artificial Intelligence and Statistics*, pp. 1232–1240, 2016.
- Jitao Sang, Yuhang Wang, Jing Zhang, Yanxu Zhu, Chao Kong, Junhong Ye, Shuyu Wei, and Jinlin Xiao. Improving weak-to-strong generalization with scalable oversight and ensemble learning. *arXiv preprint arXiv:2402.00667*, 2024.
- Tianhao Shen, Renren Jin, Yufei Huang, Chuang Liu, Weilong Dong, Zishan Guo, Xinwei Wu, Yan Liu, and Deyi Xiong. Large language model alignment: A survey. *arXiv preprint arXiv:2309.15025*, 2023.
- Rui Shu, Hung H Bui, Hirokazu Narui, and Stefano Ermon. A dirt-t approach to unsupervised domain adaptation. *arXiv preprint arXiv:1802.08735*, 2018.
- Seamus Somerstep, Felipe Maia Polo, Moulinath Banerjee, Ya’acov Ritov, Mikhail Yurochkin, and Yuekai Sun. A statistical framework for weak-to-strong generalization. *arXiv preprint arXiv:2405.16236*, 2024.
- Hwanjun Song, Minseok Kim, Dongmin Park, Yooju Shin, and Jae-Gil Lee. Learning from noisy labels with deep neural networks: A survey. *IEEE transactions on neural networks and learning systems*, 34(11):8135–8153, 2022.
- Huayi Tang and Yong Liu. Information-theoretic generalization bounds for transductive learning and its applications. *arXiv preprint arXiv:2311.04561*, 2023.
- Antti Tarvainen and Harri Valpola. Mean teachers are better role models: Weight-averaged consistency targets improve semi-supervised deep learning results. *Advances in neural information processing systems*, 30, 2017.
- Katherine Tian, Eric Mitchell, Allan Zhou, Archit Sharma, Rafael Rafailov, Huaxiu Yao, Chelsea Finn, and Christopher Manning. Just ask for calibration: Strategies for eliciting calibrated confidence scores from language models fine-tuned with human feedback. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pp. 5433–5442, 2023.
- Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*, 2023.

- Dennis Ulmer, Jes Frellsen, and Christian Hardmeier. Exploring predictive uncertainty and calibration in nlp: A study on the impact of method & data scarcity. *arXiv preprint arXiv:2210.15452*, 2022.
- Jesper E Van Engelen and Holger H Hoos. A survey on semi-supervised learning. *Machine learning*, 109(2):373–440, 2020.
- Ziqiao Wang and Yongyi Mao. Information-theoretic analysis of unsupervised domain adaptation. In *International Conference on Learning Representations*, 2023a.
- Ziqiao Wang and Yongyi Mao. Tighter information-theoretic generalization bounds from supersamples. In *International Conference on Machine Learning*, 2023b.
- David X Wu and Anant Sahai. Provable weak-to-strong generalization via benign overfitting. In *International Conference on Learning Representations*, 2025.
- Aolin Xu and Maxim Raginsky. Information-theoretic analysis of generalization capability of learning algorithms. *Advances in neural information processing systems*, 30, 2017.
- Wenkai Yang, Shiqi Shen, Guangyao Shen, Wei Yao, Yong Liu, Zhi Gong, Yankai Lin, and Ji-Rong Wen. Super (ficial)-alignment: Strong models may deceive weak models in weak-to-strong generalization. In *International Conference on Learning Representations*, 2025.
- Yuqing Yang, Yan Ma, and Pengfei Liu. Weak-to-strong reasoning. In *Findings of the Association for Computational Linguistics: EMNLP 2024*, pp. 8350–8367, 2024.
- Ruimeng Ye, Yang Xiao, and Bo Hui. Weak-to-strong generalization beyond accuracy: a pilot study in safety, toxicity, and legal reasoning. *arXiv preprint arXiv:2410.12621*, 2024.
- Zhi-Hua Zhou. A brief introduction to weakly supervised learning. *National science review*, 5(1): 44–53, 2018.
- Chiwei Zhu, Benfeng Xu, Quan Wang, Yongdong Zhang, and Zhendong Mao. On the calibration of large language models and alignment. In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pp. 9778–9795, 2023.
- Wenhong Zhu, Zhiwei He, Xiaofeng Wang, Pengfei Liu, and Rui Wang. Weak-to-strong preference optimization: Stealing reward from weak aligned model. In *International Conference on Learning Representations*, 2025.

Contents

A	Further Related Work	16
B	Main Proof	17
	B.1 Proof of Theorem 1.	17
	B.2 Extension of Theorem 1.	19
	B.3 Proof of Theorem 2.	20
	B.4 Proof of Theorem 3.	21
	B.5 Extension of Theorem 3.	23
	B.6 Proof of Theorem 4.	27
	B.7 Extension of Theorem 4.	32
C	Further Details and Results of Experiments	34
	C.1 Training Details of Experiments in Language Models	34
	C.2 Weak-to-Strong Training Procedure in Synthetic Experiments.	35

Appendix

A FURTHER RELATED WORK

Teacher-student learning paradigm. The student-teacher training paradigm (Meseguer-Brocal et al., 2019; Meng et al., 2019), which involves first training a teacher model and then using its outputs (e.g., pseudo-labels or soft targets) to guide the training of a student model, has become a cornerstone in various machine learning domains. This approach is particularly prominent in knowledge distillation (Hinton, 2015; Beyer et al., 2022), semi-supervised learning (Tarvainen & Valpola, 2017) and domain adaptation (Shu et al., 2018). It can also be used in other fields like curriculum learning (Matiisen et al., 2019) and continual learning (Lee et al., 2021). However, most prior work assumes that the teacher is either more capable or at least comparable to the student in terms of model capacity or performance. In contrast, weak-to-strong generalization (Burns et al., 2023) explores a less studied setting where the student model is significantly more capable than the teacher, which is different from the traditional assumptions of the student-teacher framework. By theoretically investigating this setting, we aim to uncover novel insights into the capabilities and limitations of weak-to-strong generalization. We hope that our theoretical investigation into it will serve as a catalyst for advancements not only in the domain of super-alignment but also in the broader landscape of teacher-student learning paradigm.

Weakly-supervised learning. Weakly-supervised learning has emerged as a powerful paradigm to address the challenges of limited labeled data by leveraging weak supervision (Ratner et al., 2020). Such weak supervision may be incomplete (i.e., only a small subset of labels are given), inexact (i.e., only coarse-grained labels are given) and inaccurate (i.e., the given labels are noisy) (Zhou, 2018). This problem setting is also closely related to label noise (Song et al., 2022) and semi-supervised learning (Van Engelen & Hoos, 2020). To address the problem of weakly-supervised learning, practical trials leverage these various forms of weak supervision, such as utilizing noisy labels (Cheng et al., 2020), coarse-grained labels (Oquab et al., 2015), and incomplete annotations (Papadopoulos et al., 2017). And most of them improving model performance within the limitations of weak supervision. In contrast, weak-to-strong generalization explores a distinct yet related direction: it investigates how a strong model, when trained on weak supervision, can not only correct the errors of the weak supervisor but also generalize to instances where the weak supervisor is uncertain or incorrect (Burns et al., 2023; Yang et al., 2025). We hope that our theoretical exploration of weak-to-strong generalization can inspire not only the field of super-alignment but also research in weakly-supervised learning.

Calibration. Calibration is an important concept about uncertainty estimation and confidence (Guo et al., 2017; Kuleshov et al., 2018; Kumar et al., 2019; Mehrtash et al., 2020) in machine learning.

There are several kinds of definition for calibration. For instance, taking expectation conditioned on the data distribution (Kull et al., 2019; Kumar et al., 2019; Roelofs et al., 2022) (Also, see Definition 3 in (Pleiss et al., 2017) and Equation (2) in (Liu et al., 2019) in the fairness literature), and (2) taking expectation conditioned on the probability score (Naeini et al., 2015; Guo et al., 2017). Researchers also investigate the calibration in natural language processing (Desai & Durrett, 2020; Guo et al., 2021; Ulmer et al., 2022; Chen et al., 2023). In recent years, the calibration of large language models has garnered significant attention (Zhu et al., 2023; Tian et al., 2023; Liang et al., 2023), with a thorough survey provided in (Geng et al., 2024). However, to the best of our knowledge, while confidence issues in weak-to-strong generalization have been investigated in (Burns et al., 2023), the role of calibration has not been sufficiently investigated. In this paper, we theoretically demonstrate how strong model’s calibration is affected in WTSG. And we believe that calibration warrants further in-depth investigation in this field.

Information-theoretic analysis. Information-theoretic analysis is commonly employed to bound the expected generalization error in supervised learning (Russo & Zou, 2016; Xu & Raginsky, 2017), with subsequent studies providing sharper bounds (Bu et al., 2020; Wang & Mao, 2023b). These bounds have been used to characterize the generalization ability of stochastic gradient-based optimization algorithms (Pensia et al., 2018). Furthermore, this theoretical framework has been extended to diverse settings, including meta-learning (Chen et al., 2021), semi-supervised learning (Aminian et al., 2022), transductive learning (Tang & Liu, 2023), and domain adaptation (Wang & Mao, 2023a). For a comprehensive overview of these developments, we refer readers to the recent monograph by Hellström et al. (2023). Nonetheless, despite its extensive application across various domains, the information-theoretic analysis of super-alignment (OpenAI, 2024), particularly in the context of weak-to-strong generalization, remains largely underexplored. In this paper, we use KL divergence to analyze weak-to-strong generalization, which is not considered in previous work. KL divergence is an information-theoretic measure between two probability distributions in information theory (Cover, 1999). And how to extend it to other information-theoretic measures remains an open question and warrants further exploration in future work.

B MAIN PROOF

B.1 PROOF OF THEOREM 1

We first state some preliminaries for the proof.

Lemma 1 (Donsker and Varadhan’s variational formula (Donsker & Varadhan, 1983)). *Let Q, P be probability measures on $\mathcal{X}$, for any bounded measurable function $f : \mathcal{X} \rightarrow \mathbb{R}$, we have*

$$D_{\text{KL}}(Q||P) = \sup_f \mathbb{E}_{x \sim Q}[f(x)] - \log \mathbb{E}_{x \sim P}[\exp f(x)].$$

Lemma 2 (Hoeffding’s lemma). *Let $X \in \mathbb{R}$ such that $a \leq X \leq b$. Then, for all $\lambda \in \mathbb{R}$,*

$$\mathbb{E} \left[e^{\lambda(X - \mathbb{E}[X])} \right] \leq \exp \left(\frac{\lambda^2(b - a)^2}{8} \right).$$

Definition 2 (Subgaussian random variable). *A random variable $X \in \mathbb{R}$ is σ -subgaussian if for any ρ ,*

$$\log \mathbb{E} \exp(\rho(X - \mathbb{E}X)) \leq \rho^2 \sigma^2 / 2.$$

Notation of probability distribution for the model output. We define the corresponding probability distributions for prediction of F_{sw} and F_w . Recall that for $F_w, F_{sw} : \mathcal{X} \rightarrow \mathcal{Y}$ and $x \in \mathcal{X}$:

$$d_{\mathcal{P}}(F_w, F_{sw}) = \mathbb{E}_x \left[\sum_{j=1}^k [F_w(x)]_j \log \frac{[F_w(x)]_j}{[F_{sw}(x)]_j} \right] = \mathbb{E}_x [D_{\text{KL}}(F_w(x), F_{sw}(x))],$$

where D_{KL} is the discrete version of KL divergence. $\forall x \in \mathcal{X}$, we know that $\sum_{j=1}^k [F_w(x)]_j = 1$. Therefore, given the class space $C_k = \{1, \dots, k\}$, we define a probability distribution $\mathcal{P}_w(x)$ with the probability density function p_w , where $j \in C_k$ and

$$p_w(j) = [F_w(x)]_j. \quad (13)$$

Using this method, we also define the probability distribution $\mathcal{P}_{sw}(x)$ for $F_w(x)$.

Now we start the proof.

Proof. For better readability, we divide the proof into several steps.

The first step. Given the probability distributions $\mathcal{P}_w(x)$ and $\mathcal{P}_{sw}(x)$ above, the first step is motivated by Lemma A.2 from (Wang & Mao, 2023a). For any $x \in \mathcal{X}$, $j \in C_k$, $g : C_k \rightarrow \mathbb{R}$ and assume that g is σ -subgaussian (we will specify σ later). Let $f = t \cdot g$ for any $t \in \mathbb{R}$. We have

$$\begin{aligned}
D_{\text{KL}}(F_w(x) \| F_{sw}(x)) &= D_{\text{KL}}(\mathcal{P}_w(x) \| \mathcal{P}_{sw}(x)) \\
&= \sup_t \mathbb{E}_{j' \sim \mathcal{P}_w(x)} [t \cdot g(j')] - \log \mathbb{E}_{j \sim \mathcal{P}_{sw}(x)} [\exp(t \cdot g(j))] \quad (\text{Lemma 1}) \\
&= \sup_t \mathbb{E}_{j' \sim \mathcal{P}_w(x)} [tg(j')] - \log \mathbb{E}_{j \sim \mathcal{P}_{sw}(x)} [\exp t (g(j) - \mathbb{E}_{j \sim \mathcal{P}_{sw}(x)}[g(j)] + \mathbb{E}_{j \sim \mathcal{P}_{sw}(x)}[g(j)])] \\
&= \sup_t \mathbb{E}_{j' \sim \mathcal{P}_w(x)} [tg(j')] - \mathbb{E}_{j \sim \mathcal{P}_{sw}(x)} [tg(j)] - \log \mathbb{E}_{j \sim \mathcal{P}_{sw}(x)} [\exp t (g(j) - \mathbb{E}_{j \sim \mathcal{P}_{sw}(x)}[g(j)])] \\
&\geq \sup_t t (\mathbb{E}_{j' \sim \mathcal{P}_w(x)} [g(j')] - \mathbb{E}_{j \sim \mathcal{P}_{sw}(x)} [g(j)]) - t^2 \sigma^2 / 2.
\end{aligned}$$

(Subgaussianity)

The second step. The second step is associating the above result with $d_{\mathcal{P}}(F_w, F_{sw})$. In particular, by taking expectations of x on both sides of the above inequality, we obtain

$$d_{\mathcal{P}}(F_w, F_{sw}) = \mathbb{E}_x D_{\text{KL}}(F_w(x) \| F_{sw}(x)) \geq \sup_t \underbrace{t (\mathbb{E}_x \mathbb{E}_{j' \sim \mathcal{P}_w(x)} [g(j')] - \mathbb{E}_x \mathbb{E}_{j \sim \mathcal{P}_{sw}(x)} [g(j)])}_{\phi(t)} - t^2 \sigma^2 / 2.$$

Note that $\phi(t)$ is a quadratic function of t . Therefore, by AM–GM inequality, we find the maximum of this quadratic function:

$$\phi(t) \leq \frac{1}{2\sigma^2} (\mathbb{E}_x \mathbb{E}_{j' \sim \mathcal{P}_w(x)} [g(j')] - \mathbb{E}_x \mathbb{E}_{j \sim \mathcal{P}_{sw}(x)} [g(j)])^2 = \sup_t \phi(t) \leq d_{\mathcal{P}}(F_w, F_{sw}).$$

Subsequently, there holds

$$|\mathbb{E}_x \mathbb{E}_{j' \sim \mathcal{P}_w(x)} [g(j')] - \mathbb{E}_x \mathbb{E}_{j \sim \mathcal{P}_{sw}(x)} [g(j)]| \leq \sqrt{2\sigma^2 d_{\mathcal{P}}(F_w, F_{sw})}. \quad (14)$$

The third step. The third step is constructing g to associate the above result with $d_{\mathcal{P}}(F^*, F_{sw})$ and $d_{\mathcal{P}}(F^*, F_w)$. Specifically, given a probability distribution $\mathcal{P}_g$ with the density function p_g , we define function $g : C_k \rightarrow (0, 1]$ associated with $\mathcal{P}_g$:

$$g(j) \triangleq \frac{[F^*(x)]_j}{p_g(j)} \log \frac{[F^*(x)]_j}{p_g(j)}, \quad \text{for } j \in C_k.$$

We have

$$\begin{aligned}
\mathbb{E}_x \mathbb{E}_{j \sim \mathcal{P}_g} [g(j)] &= \mathbb{E}_x \mathbb{E}_{j \sim \mathcal{P}_g} \left[\frac{[F^*(x)]_j}{p_g(j)} \log \frac{[F^*(x)]_j}{p_g(j)} \right] \\
&= \mathbb{E}_x \left[\sum_{j \in C_k} p_g(j) \cdot \frac{[F^*(x)]_j}{p_g(j)} \cdot \log \frac{[F^*(x)]_j}{p_g(j)} \right] \\
&= \mathbb{E}_x \left[\sum_{j \in C_k} [F^*(x)]_j \cdot \log \frac{[F^*(x)]_j}{p_g(j)} \right]
\end{aligned}$$

Recall the definition of $\mathcal{P}_{sw}$ and $\mathcal{P}_w$ in (13), we replace $\mathcal{P}_g$ with $\mathcal{P}_{sw}$ and $\mathcal{P}_w$ in the above equation:

$$\mathbb{E}_x \mathbb{E}_{j' \sim \mathcal{P}_{sw}} [g(j')] = \mathbb{E}_x \left[\sum_{j=1} [F^*(x)]_j \log \frac{[F^*(x)]_j}{[F_{sw}(x)]_j} \right] = d_{\mathcal{P}}(F^*, F_{sw}),$$

$$\mathbb{E}_x \mathbb{E}_{j \sim \mathcal{P}_w} [g(j)] = \mathbb{E}_x \left[\sum_{j=1} [F^*(x)]_j \log \frac{[F^*(x)]_j}{[F_w(x)]_j} \right] = d_{\mathcal{P}}(F^*, F_w).$$

Substitute the above into (14):

$$|d_{\mathcal{P}}(F^*, F_{sw}) - d_{\mathcal{P}}(F^*, F_w)| \leq \sqrt{2\sigma^2 d_{\mathcal{P}}(F_w, F_{sw})}. \quad (15)$$

The final step. Finally, we obtain the subgaussian factor R of function g by using the fact that g is bounded. For simplicity, we use Hoeffding's Lemma (Lemma 2) to obtain the subgaussian factor R . However, it can be more precisely determined using advanced techniques in learning theory literature (for instance, see Remark 2.14 in (Li et al., 2024), where $\alpha = 2$ recovers the subgaussian setting).

Recall that the output domain $\mathcal{Y} \subseteq \mathbb{R}^k$, where $\forall y = (y_1, \dots, y_k)^T \in \mathcal{Y}$, there holds $\sum_{i=1}^k y_i = 1$ and $0 < y_i \leq 1$. In other words, $\exists \gamma > 0$ such that $0 < \gamma \leq y_i \leq 1$. It means that $g(j) \in [-\frac{1}{\gamma} \log \frac{1}{\gamma}, \frac{1}{\gamma} \log \frac{1}{\gamma}]$. According to Lemma 2, $\forall \lambda \in \mathbb{R}$, we have

$$\mathbb{E} \left[e^{\lambda(g(j) - \mathbb{E}[g(j)])} \right] \leq \exp \left(\frac{\lambda^2 \left(\frac{1}{\gamma} \log \frac{1}{\gamma} \right)^2}{2} \right).$$

In other words, $g(j)$ is σ -subgaussian, where $\sigma = \frac{1}{\gamma} \log \frac{1}{\gamma}$. Substitute it into Equation (15) and we obtain:

$$|d_{\mathcal{P}}(F^*, F_{sw}) - d_{\mathcal{P}}(F^*, F_w)| \leq C_1 \sqrt{d_{\mathcal{P}}(F_w, F_{sw})},$$

where the constant $C_1 = \frac{\sqrt{2}}{\gamma} \log \frac{1}{\gamma}$. The proof is complete. $\square$

B.2 EXTENSION OF THEOREM 1

In this section, we extend Theorem 1 to output distribution divergence in regression.

Proof. Denote $d_{\text{KL}}(f \| g) = D_{\text{KL}}(\mathcal{P}_f \| \mathcal{P}_g) = \int_{\mathcal{X}} f(x) \log \frac{f(x)}{g(x)} dx$. Let $(\mathcal{X}, \mathcal{F}, \mathcal{P}_{sw})$, $(\mathcal{X}, \mathcal{F}, \mathcal{P}_w)$ be two probability spaces. Denoting F_{sw} and F_w the densities of the measures. Therefore,

$$\int_{\mathcal{X}} F_{sw}(x) dx = \int_{\mathcal{X}} F_w(x) dx = 1.$$

Let $x \in \mathcal{X}$, $g : \mathcal{X} \rightarrow \mathbb{R}$ and assume that g is R -subgaussian (we will specify R later). Let $f = t \cdot g$ for any $t \in \mathbb{R}$. By Lemma 1, we have

$$\begin{aligned} d_{\text{KL}}(F_w \| F_{sw}) &= D_{\text{KL}}(\mathcal{P}_w \| \mathcal{P}_{sw}) \\ &= \sup_t \mathbb{E}_{x' \sim \mathcal{P}_w} [t \cdot g(x')] - \log \mathbb{E}_{x \sim \mathcal{P}_{sw}} [\exp(t \cdot g(x))] \\ &= \sup_t \mathbb{E}_{x' \sim \mathcal{P}_w} [tg(x')] - \log \mathbb{E}_{x \sim \mathcal{P}_{sw}} [\exp t(g(x) - \mathbb{E}_{x \sim \mathcal{P}_{sw}}[g(x)] + \mathbb{E}_{x \sim \mathcal{P}_{sw}}[g(x)])] \\ &= \sup_t \mathbb{E}_{x' \sim \mathcal{P}_w} [tg(x')] - \mathbb{E}_{x \sim \mathcal{P}_{sw}} [tg(x)] - \log \mathbb{E}_{x \sim \mathcal{P}_{sw}} [\exp t(g(x) - \mathbb{E}_{x \sim \mathcal{P}_{sw}}[g(x)])] \\ &\geq \sup_t t \underbrace{(\mathbb{E}_{x' \sim \mathcal{P}_w} [g(x')] - \mathbb{E}_{x \sim \mathcal{P}_{sw}} [g(x)])}_{\phi(t)} - t^2 R^2 / 2. \end{aligned} \quad (\text{Subgaussianity})$$

Let

$$\phi(t) = t(\mathbb{E}_{x' \sim \mathcal{P}_w} [g(x')] - \mathbb{E}_{x \sim \mathcal{P}_{sw}} [g(x)]) - t^2 R^2 / 2,$$

which is a quadratic function of t . Therefore,

$$\phi(t) \leq \frac{1}{2R^2} (\mathbb{E}_{x' \sim \mathcal{P}_{sw}} [g(x')] - \mathbb{E}_{x \sim \mathcal{P}_w} [g(x)])^2 = \sup_t \phi(t) \leq D_{\text{KL}}(\mathcal{P}_w \| \mathcal{P}_{sw}) = d_{\text{KL}}(F_w \| F_{sw}).$$

So we have

$$|\mathbb{E}_{x' \sim \mathcal{P}_{sw}} [g(x')] - \mathbb{E}_{x \sim \mathcal{P}_w} [g(x)]| \leq \sqrt{2R^2 d_{\text{KL}}(F_w \| F_{sw})}. \quad (16)$$

Given a probability space $(\mathcal{X}, \mathcal{F}, \mathcal{P}_g)$ with the density function p_g . Define

$$g(x) \triangleq \frac{F^*(x)}{p_g(x)} \log \frac{F^*(x)}{p_g(x)}.$$

So there holds

$$\mathbb{E}_{x \sim \mathcal{P}_g} [g(x)] = \int_{\mathcal{X}} F^*(x) \cdot \log \frac{F^*(x)}{p_g(x)} dx.$$

Replace $\mathcal{P}_g$ with $\mathcal{P}_{sw}$ and $\mathcal{P}_w$:

$$\begin{aligned} \mathbb{E}_{x' \sim \mathcal{P}_{sw}} [g(x')] &= \int_{\mathcal{X}} F^*(x) \log \frac{F^*(x)}{F_{sw}(x)} dx = d_{\mathcal{P}}(F^*, F_{sw}), \\ \mathbb{E}_{x \sim \mathcal{P}_w} [g(x)] &= \int_{\mathcal{X}} F^*(x) \log \frac{F^*(x)}{F_w(x)} dx = d_{\mathcal{P}}(F^*, F_w). \end{aligned}$$

Substitute them back into (16) and we obtain:

$$|d_{\mathcal{P}}(F^*, F_{sw}) - d_{\mathcal{P}}(F^*, F_w)| \leq \sqrt{2R^2 d_{\text{KL}}(F_w \| F_{sw})}. \quad (17)$$

Finally, recall that the output domain $\mathcal{Y} = \{y \in \mathbb{R} | 0 < y \leq 1\}$. In other words, $\exists \gamma > 0$ such that $\mathcal{Y} = \{y \in \mathbb{R} | 0 < \gamma \leq y \leq 1\}$. It means that $g(x) \in [-\frac{1}{\gamma} \log \frac{1}{\gamma}, \frac{1}{\gamma} \log \frac{1}{\gamma}]$. According to Lemma 2, $\forall \lambda \in \mathbb{R}$, we have

$$\mathbb{E} \left[e^{\lambda(g(x) - \mathbb{E}[g(x)])} \right] \leq \exp \left(\frac{\lambda^2 \left(\frac{1}{\gamma} \log \frac{1}{\gamma} \right)^2}{2} \right).$$

In other words, $g(x)$ is R -subgaussian, where $R = \frac{1}{\gamma} \log \frac{1}{\gamma}$. Substitute it into Equation (17) and we obtain:

$$|d_{\mathcal{P}}(F^*, F_{sw}) - d_{\mathcal{P}}(F^*, F_w)| \leq \sqrt{C_1 d_{\mathcal{P}}(F_w, F_{sw})},$$

where the constant $C_1 = 2 \left(\frac{1}{\gamma} \log \frac{1}{\gamma} \right)^2$. □

B.3 PROOF OF THEOREM 2

Total variation distance is introduced for our proof.

Definition 3 (Total Variation Distance). *Given two probability distributions P and Q , the Total Variation (TV) distance between P and Q is*

$$D_{\text{TV}}(P \| Q) = \frac{1}{2} \int_{x \in \mathcal{X}} |P(x) - Q(x)| dx.$$

Note that $D_{\text{TV}}(P \| Q) \in [0, 1]$. Also, $D_{\text{TV}}(P \| Q) = 0$ if and only if P and Q coincides, and $D_{\text{TV}}(P \| Q) = 1$ if and only if P and Q are disjoint.

Let the calibrated bayes score function $F^b : \mathcal{X} \rightarrow \mathcal{Y}$ that satisfies $\forall i \in \{1, \dots, k\}$, $[F^b(x)]_i = \mathbb{P}(Y_i = 1 | X = x)$, where $Y = [Y_1, \dots, Y_k]^T \in \{0, 1\}^k$ and $\|Y\|_1 = 1$. Now we start our proof.

Proof. Consider the definition of MCE in Equation (9). Notice that

$$MCE(F) = \sum_{i=1}^k \mathbb{E}_X |[F(X)]_i - \mathbb{P}[Y_i = 1 | F(X)]_i|$$

$$\begin{aligned}
&= \mathbb{E}_X \left[\sum_{i=1}^k |[F(X)]_i - \mathbb{P}[Y_i = 1 | [F(X)]_i]| \right] \\
&= \mathbb{E}_X \|F(X) - F^b(X)\|_1,
\end{aligned}$$

and $MCE(F) \in [0, 2]$.

So there holds

$$\begin{aligned}
MCE(F_w) - MCE(F_{sw}) &= \mathbb{E}_X \|F_w(X) - F^b(X)\|_1 - \mathbb{E}_X \|F_{sw}(X) - F^b(X)\|_1 \\
&\leq \mathbb{E}_X \|F_w(X) - F_{sw}(X)\|_1 && \text{(Triangle inequality)} \\
&= 2 \cdot \mathbb{E}_X D_{TV}(F_w(X), F_{sw}(X)) \\
&\leq 2 \cdot \mathbb{E}_X \sqrt{1 - \exp(-D_{KL}(F_w(X), F_{sw}(X)))} && \text{(Bretagnolle–Huber inequality)} \\
&\leq 2 \cdot \sqrt{1 - \exp(-\mathbb{E}_X D_{KL}(F_w(X), F_{sw}(X)))} && \text{(Jensen's inequality)} \\
&= 2 \cdot \sqrt{1 - \exp(-d_{\mathcal{P}}(F_w, F_{sw}))}. && \text{(Definition of } d_{\mathcal{P}} \text{ for KL divergence loss)}
\end{aligned}$$

Likewise,

$$\begin{aligned}
MCE(F_{sw}) - MCE(F_w) &= \mathbb{E}_X \|F_{sw}(X) - F^b(X)\|_1 - \mathbb{E}_X \|F_w(X) - F^b(X)\|_1 \\
&\leq \mathbb{E}_X \|F_{sw}(X) - F_w(X)\|_1 && \text{(Triangle inequality)} \\
&\leq \mathbb{E}_X \|F_w(X) - F_{sw}(X)\|_1 && \text{(Symmetry)} \\
&= 2 \cdot \mathbb{E}_X D_{TV}(F_w(X), F_{sw}(X)) \\
&\leq 2 \cdot \sqrt{1 - \exp(-d_{\mathcal{P}}(F_w, F_{sw}))}. && \text{(Using the derivation above)}
\end{aligned}$$

Combining the above, we have

$$\begin{aligned}
MCE(F_{sw}) - MCE(F_w) &\leq 2 \cdot \sqrt{1 - \exp(-d_{\mathcal{P}}(F_w, F_{sw}))}, \\
MCE(F_w) - MCE(F_{sw}) &\leq 2 \cdot \sqrt{1 - \exp(-d_{\mathcal{P}}(F_w, F_{sw}))}.
\end{aligned}$$

The proof is complete. □

B.4 PROOF OF THEOREM 3

We first restate a lemma for our proof. Recall that the strong model learns from a linear function class $\mathcal{F} : \mathbb{R}^{d_s} \rightarrow \mathbb{R}$ of fine-tuning tasks. Recall also that we denote the strong model representation map by $h_s : \mathbb{R}^d \rightarrow \mathbb{R}^{d_s}$. Let $V_s = \{f \circ h_s : f \in \mathcal{F}\}$ be the set of all tasks in $\mathcal{F}$ composed with the strong model representation. We first observe that V_s is also a convex set.

Lemma 3 (Charikar et al. (2024)). *V_s is a convex set.*

Proof. Fix $f, g \in \mathcal{F}$, and consider $f \circ h_s, g \circ h_s \in V_s$. Fix any $\lambda \in [0, 1]$. Since $\mathcal{F}$ is the linear function class so that it is a convex set, there exists $p \in \mathcal{F}$ such that for all $y \in \mathbb{R}^{d_s}$, $p(y) = \lambda f(y) + (1 - \lambda)g(y)$. Now, fix any $x \in \mathbb{R}^d$. Then, we have that

$$\begin{aligned}
\lambda(f \circ h_s)(x) + (1 - \lambda)(g \circ h_s)(x) &= \lambda f(h_s(x)) + (1 - \lambda)g(h_s(x)) = p(h_s(x)) = (p \circ h_s)(x), \\
&\text{and hence } \lambda(f \circ h_s) + (1 - \lambda)(g \circ h_s) = p \circ h_s \in V_s. \quad \square
\end{aligned}$$

Now we start the proof.

Proof. For any $f, g \in \mathcal{X} \rightarrow \mathcal{Y}$, denote $d_{KL}(f \| g) = \int_{\mathcal{X}} f(x) \log \frac{f(x)}{g(x)} dx$.

Given any $g \in V_s$, observe that

$$\begin{aligned}
d_{\text{KL}}(g\|F_w) &= \int_{\mathcal{X}} g(x) \log \frac{g(x)}{F_w(x)} dx \\
&= \int_{\mathcal{X}} g(x) \log \left(\frac{g(x)}{F_{sw}(x)} \cdot \frac{F_{sw}(x)}{F_w(x)} \right) dx \\
&= d_{\text{KL}}(g\|F_{sw}) + d_{\text{KL}}(F_{sw}\|F_w) - d_{\text{KL}}(F_{sw}\|F_w) + \int_{\mathcal{X}} g(x) \log \frac{F_{sw}(x)}{F_w(x)} dx \\
&= d_{\text{KL}}(g\|F_{sw}) + d_{\text{KL}}(F_{sw}\|F_w) + \underbrace{\int_{\mathcal{X}} (g(x) - F_{sw}(x)) \log \frac{F_{sw}(x)}{F_w(x)} dx}_{Q_1}. \quad (18)
\end{aligned}$$

Now our goal is to judge whether $Q_1 \geq 0$.

Recall that

$$f_{sw} = \operatorname{argmin}_f d_{\text{KL}}(f \circ h_s\|F_w).$$

In other words, F_{sw} is the *projection* of F_w onto the convex set V_s . Therefore:

$$d_{\text{KL}}(g\|F_w) \geq d_{\text{KL}}(F_{sw}\|F_w).$$

Therefore,

$$d_{\text{KL}}(g\|F_{sw}) + Q_1 \geq 0. \quad (19)$$

Now, fix $t \in (0, 1)$, and consider the function $g = F_{sw} + t(F^* - F_{sw})$.

$$\begin{aligned}
d_{\text{KL}}(g\|F_{sw}) &= \int_{\mathcal{X}} g(x) \log \frac{g(x)}{F_{sw}(x)} dx \\
&= - \int_{\mathcal{X}} g(x) \log \frac{F_{sw}(x)}{g(x)} dx \\
&= - \int_{\mathcal{X}} g(x) \log \left[1 - \frac{t(F^*(x) - F_{sw}(x))}{g(x)} \right] dx \\
&= \int_{\mathcal{X}} g(x) \left[\frac{t(F^*(x) - F_{sw}(x))}{g(x)} + \mathcal{O}(t^2) \right] dx \\
&= \int_{\mathcal{X}} [t(F^*(x) - F_{sw}(x)) + \mathcal{O}(t^2)] dx \quad (\text{Taylor expansion}) \\
&= \mathcal{O}(t^2). \quad (\int_{\mathcal{X}} F_{sw}(x) dx = \int_{\mathcal{X}} F^*(x) dx = 1)
\end{aligned}$$

While

$$\begin{aligned}
Q_1 &= \int_{\mathcal{X}} (g(x) - F_{sw}(x)) \log \frac{F_{sw}(x)}{F_w(x)} dx \\
&= t \cdot \int_{\mathcal{X}} (F^*(x) - F_{sw}(x)) \log \frac{F_{sw}(x)}{F_w(x)} dx \\
&= \mathcal{O}(t).
\end{aligned}$$

Recall Equation (19) that

$$\underbrace{d_{\text{KL}}(F_{sw}\|g)}_{\mathcal{O}(t^2)} + \underbrace{Q_1}_{\mathcal{O}(t)} \geq 0,$$

which means $Q_1 \geq 0$. So we have

$$\int_{\mathcal{X}} (F^*(x) - F_{sw}(x)) \log \frac{F_{sw}(x)}{F_w(x)} dx \geq 0.$$

Let $g = F^*$ in Equation (18) and we can prove the result $d_{\mathcal{P}}(F^*, F_{sw}) \leq d_{\mathcal{P}}(F^*, F_w) - d_{\mathcal{P}}(F_{sw}, F_w)$.

□

B.5 EXTENSION OF THEOREM 3

We first introduce some definitions.

Definition 4 (Itakura–Saito Divergence (Itakura, 1968; Févotte et al., 2009; Prasetyawan & Takamichi, 2020)). *Given two probability distributions P and Q , the Itakura–Saito divergence between them is defined as*

$$D_{\text{IS}}(P\|Q) = \int_{\mathcal{X}} \left(\frac{P(x)}{Q(x)} - \log \frac{P(x)}{Q(x)} - 1 \right) dx.$$

Similar to the KL divergence, the Itakura–Saito divergence is also a Bregman divergence (Dhillon, 2007).

Definition 5 (Weighted Itakura–Saito Divergence (Chu & Messerschmitt, 1982)). *Given two probability distributions P and Q , the weighted Itakura–Saito divergence between them is defined as*

$$D_{\text{WIS}}(P\|Q) = \int_{\mathcal{X}} w(x) \left(\frac{P(x)}{Q(x)} - \log \frac{P(x)}{Q(x)} - 1 \right) dx,$$

where $w(x)$ is the weight function.

We also define the inner product of functions

$$\langle f, g \rangle \triangleq \int_{\mathcal{X}} f(x)g(x)dx.$$

Now we present the theoretical extension of Theorem 3 to forward KL divergence loss in WTSG.

Corollary 1. *Given F^* , F_w and F_{sw} defined above. Let $d_{\mathcal{P}}$ be the output distribution divergence and consider WTSG in Equation (6) using forward KL divergence loss. Assume that $\exists f_s \in \mathcal{F}_s$ such that $F_s = F^*$. Consider $\mathcal{F}_s$ that satisfies Assumption 1. If the weighted Itakura–Saito divergence $D_{\text{WIS}}(F^*\|F_{sw}) \leq 0$ with the weight function $w = F_{sw} - F_w$, then:*

$$d_{\mathcal{P}}(F_{sw}, F^*) \leq d_{\mathcal{P}}(F_w, F^*) - d_{\mathcal{P}}(F_w, F_{sw}). \quad (20)$$

Corollary 1 shows that Charikar et al. (2024) can be extended to our setting if we introduce another assumption, which comes from the technical challenges of theoretically analyzing the non-linear nature of KL divergence. Denote $F^+ = \frac{F^*}{F_{sw}} - 1 - \log \frac{F^*}{F_{sw}}$, then the assumption $D_{\text{WIS}}(F^*\|F_{sw}) \leq 0$ is equivalent to $\langle F_{sw} - F_w, F^+ \rangle \leq 0$. Note that $\forall x \in \mathcal{X}$, $F_+(x) \geq 0$ always holds, and a very small or large value of $\frac{F^*(x)}{F_{sw}(x)}$ generally contributes to a large $F_+(x)$. To make $\langle F_{sw} - F_w, F^+ \rangle \leq 0$ more likely to hold, we expect $F_{sw}(x) \leq F_w(x)$ if $F^*(x)$ is small or large. In general, since this condition cannot be guaranteed to hold universally, the inequality in Equation (11) may fail to hold. This reveals a key discrepancy between the square function (as considered in Charikar et al. (2024)) and the KL divergence (in this work) within the WTSG framework—a phenomenon that will be empirically validated through our experiments.

Proof sketch of Charikar et al. (2024). For the proof technique, Charikar et al. (2024) constructs a function within a convex set. By exploiting the property of projection and square function, they demonstrate that $\mathcal{O}(t) + \mathcal{O}(t^2)$ is non-negative as $t \rightarrow 0^+$. Consequently, the first-order term must be non-negative, which proves the result.

Proof sketch of ours. Extending the proof framework from Theorem 1 in (Charikar et al., 2024) presents several challenges. First, due to the properties of KL divergence, the constructed function does not lie within the convex set. To address this issue, we employ a first-order Taylor expansion and introduce a remainder term. Secondly, because of the remainder, we derive that $\mathcal{O}(t) + \mathcal{O}(t) + \mathcal{O}(t^2)$ is non-negative. Consequently, we must assume that one of the first-order terms is non-positive to ensure that the other first-order term is non-negative, which allows us to prove the result. However, if the first-order term is positive, the second first-order term might also remain non-negative.

Now we start our proof of Corollary 1. Some Taylor expansion claims used in the proof (Claim 1, Claim 2 and Claim 3) are provided at the end of the proof.

Proof. For any $f, g \in \mathcal{X} \rightarrow \mathcal{Y}$, denote $d_{\text{KL}}(f\|g) = \int_{\mathcal{X}} f(x) \log \frac{f(x)}{g(x)} dx$.

Given any $g \in V_s$, observe that

$$\begin{aligned}
d_{\text{KL}}(F_w\|g) &= \int_{\mathcal{X}} F_w(x) \log \frac{F_w(x)}{g(x)} dx \\
&= \int_{\mathcal{X}} F_w(x) \log \left(\frac{F_w(x)}{F_{sw}(x)} \cdot \frac{F_{sw}(x)}{g(x)} \right) dx \\
&= \int_{\mathcal{X}} F_w(x) \log \frac{F_w(x)}{F_{sw}(x)} dx + \int_{\mathcal{X}} F_w(x) \log \frac{F_{sw}(x)}{g(x)} dx \\
&= d_{\text{KL}}(F_w\|F_{sw}) + d_{\text{KL}}(F_{sw}\|g) - d_{\text{KL}}(F_{sw}\|g) + \int_{\mathcal{X}} F_w(x) \log \frac{F_{sw}(x)}{g(x)} dx \\
&= d_{\text{KL}}(F_w\|F_{sw}) + d_{\text{KL}}(F_{sw}\|g) + \underbrace{\int_{\mathcal{X}} (F_w(x) - F_{sw}(x)) \log \frac{F_{sw}(x)}{g(x)} dx}_{Q_1}. \quad (21)
\end{aligned}$$

Now our goal is to judge whether $Q_1 \geq 0$.

Recall that

$$f_{sw} = \operatorname{argmin}_f d_{\text{KL}}(F_w\|f \circ h_s).$$

In other words, F_{sw} is the *projection* of F_w onto the convex set V_s . Therefore:

$$d_{\text{KL}}(F_w\|g) \geq d_{\text{KL}}(F_w\|F_{sw}).$$

And hence

$$\begin{aligned}
&d_{\text{KL}}(F_w\|g) - d_{\text{KL}}(F_w\|F_{sw}) \geq 0 \\
&\Rightarrow \int_{\mathcal{X}} F_w(x) \log \frac{F_w(x)}{g(x)} dx - \int_{\mathcal{X}} F_w(x) \log \frac{F_w(x)}{F_{sw}(x)} dx \geq 0 \\
&\Rightarrow \underbrace{\int_{\mathcal{X}} F_w(x) \log \frac{F_{sw}(x)}{g(x)} dx}_{Q_2} \geq 0.
\end{aligned}$$

Therefore,

$$Q_2 = d_{\text{KL}}(F_{sw}\|g) + Q_1 \geq 0. \quad (22)$$

Now, fix $t \in (0, 1)$, and consider the functions

$$w(x) = (F_{sw}(x)) \cdot \left(\frac{F^*(x)}{F_{sw}(x)} \right)^t, \quad (23)$$

$$w'(x) = F_{sw}(x) + t(F^*(x) - F_{sw}(x)). \quad (24)$$

It is clear that $w(x) > 0$, $w'(x) > 0$. And according to Claim 1, we have

$$w(x) = F_{sw}(x) \cdot \left[1 + t \log \frac{F^*(x)}{F_{sw}(x)} + \frac{1}{2} t^2 \left(\log \frac{F^*(x)}{F_{sw}(x)} \right)^2 \left(\frac{F^*(x)}{F_{sw}(x)} \right)^\xi \right],$$

where $\xi \in (0, t)$. It means that

$$\begin{aligned}
w'(x) - w(x) &= t \cdot F_{sw}(x) \underbrace{\left(\frac{F^*(x)}{F_{sw}(x)} - 1 - \log \frac{F^*(x)}{F_{sw}(x)} \right)}_{\mathcal{O}(t)} - t^2 \cdot \frac{F_{sw}(x)}{2} \underbrace{\left(\log \frac{F^*(x)}{F_{sw}(x)} \right)^2 \left(\frac{F^*(x)}{F_{sw}(x)} \right)^\xi}_{\mathcal{O}(t^2)}. \quad (25)
\end{aligned}$$

And $\frac{F^*(x)}{F_{sw}(x)} - 1 - \log \frac{F^*(x)}{F_{sw}(x)} > 0$. It means that as $t \rightarrow 0^+$, $w'(x) - w(x) > 0$ and $w'(x) - w(x) = \mathcal{O}(t)$.

Define

$$\begin{aligned}
\phi(x) &\triangleq \log \frac{F_{sw}(x)}{w'(x)} - \log \frac{F_{sw}(x)}{w(x)} \\
&= \log \frac{w(x)}{w'(x)} \\
&= \log \left(1 - \frac{w'(x) - w(x)}{w'(x)} \right) \\
&= - \underbrace{\frac{w'(x) - w(x)}{w'(x)}}_{\mathcal{O}(t)} - \frac{1}{2} \underbrace{\left(\frac{w'(\zeta) - w(\zeta)}{w'(\zeta)} \right)^2}_{\mathcal{O}(t^2)}, \tag{Claim 2}
\end{aligned}$$

where ζ is between 0 and $\frac{w'(x) - w(x)}{w'(x)}$. As $t \rightarrow 0^+$, we know that $\phi(x) \leq 0$ and $\phi(x) = \mathcal{O}(t)$.

Combining (23) with (26) and we have:

$$\begin{aligned}
\log \frac{F_{sw}(x)}{w(x)} &= t \log \frac{F_{sw}(x)}{F^*(x)}, \\
\text{and } \log \frac{F_{sw}(x)}{w'(x)} &= t \log \frac{F_{sw}(x)}{F^*(x)} + \phi(x). \tag{27}
\end{aligned}$$

Note that $F_{sw} \in V_s$, which is a convex set (Lemma 3). Also, $F^* \in V_s$ (Realizability). Therefore, according to the definition of convexity, we have $w' \in V_s$. Hence, substituting w' for g in (22) and consider Equation (27), we get

$$\begin{aligned}
Q_2 &= \underbrace{\int_{\mathcal{X}} F_{sw}(x) \log \frac{F_{sw}(x)}{g(x)} dx}_{d_{\text{KL}}(F_{sw} \| g)} + \underbrace{\int_{\mathcal{X}} (F_w(x) - F_{sw}(x)) \log \frac{F_{sw}(x)}{g(x)} dx}_{Q_1} \geq 0, \\
\Rightarrow d_{\text{KL}}(F_{sw} \| w') &+ \int_{\mathcal{X}} (F_w(x) - F_{sw}(x)) \cdot \phi(x) dx + \\
&\quad t \int_{\mathcal{X}} (F_w(x) - F_{sw}(x)) \log \frac{F_{sw}(x)}{F^*(x)} dx \geq 0. \tag{28}
\end{aligned}$$

Here, we address these three components individually. Our goal is to show that the first term is $\mathcal{O}(t^2)$, while the second term, which is $\mathcal{O}(t)$, is negative. Considering these collectively, the third term, also $\mathcal{O}(t)$, must be positive.

The first term. Denote $u(x) = F^*(x) - F_{sw}(x)$. Note that

$$\begin{aligned}
d_{\text{KL}}(F_{sw} \| w') &= \int_{\mathcal{X}} (F_{sw}(x)) \log \frac{F_{sw}(x)}{F_{sw}(x) + t \cdot u(x)} dx \\
&= \int_{\mathcal{X}} (F_{sw}(x)) \log \left(1 - \frac{t \cdot u(x)}{F_{sw}(x) + t \cdot u(x)} \right) dx \\
&= -t \int_{\mathcal{X}} \frac{F_{sw}(x) \cdot u(x)}{F_{sw}(x) + t \cdot u(x)} dx - \frac{1}{2} t^2 \int_{\mathcal{X}} \frac{F_{sw}(x) \cdot (u(\zeta'))^2}{(w'(\zeta'))^2} dx \tag{Claim 2} \\
&= -t \int_{\mathcal{X}} \left[u(x) - \frac{t \cdot (u(x))^2}{F_{sw}(x) + t \cdot u(x)} \right] dx - \frac{1}{2} t^2 \int_{\mathcal{X}} \frac{F_{sw}(x) \cdot (u(\zeta'))^2}{(w'(\zeta'))^2} dx \\
&= -t \underbrace{\int_{\mathcal{X}} u(x) dx}_0 + t^2 \int_{\mathcal{X}} \frac{(u(x))^2}{F_{sw}(x) + t \cdot u(x)} dx - t^2 \int_{\mathcal{X}} \frac{F_{sw}(x) \cdot (u(\zeta'))^2}{2(w'(\zeta'))^2} dx \\
&= t^2 \int_{\mathcal{X}} \left[\frac{(u(x))^2}{F_{sw}(x) + t \cdot u(x)} - \frac{F_{sw}(x) \cdot (u(\zeta'))^2}{2(w'(\zeta'))^2} \right] dx
\end{aligned}$$

where ζ' is between 0 and $\frac{t \cdot u(x)}{F_{sw}(x) + t \cdot u(x)}$. Therefore, taking the limit as $t \rightarrow 0^+$, we get that $d_{\text{KL}}(F_{sw} \| w') = \mathcal{O}(t^2)$.

The second term. Recall Equation (25) that

$$\begin{aligned} w'(x) - w(x) &= t \cdot F_{sw}(x) \left(\frac{F^*(x)}{F_{sw}(x)} - 1 - \log \frac{F^*(x)}{F_{sw}(x)} \right) + \mathcal{O}(t^2), \\ w'(x) &= F_{sw}(x) \left(1 + t \left(\frac{F^*(x)}{F_{sw}(x)} - 1 \right) \right). \end{aligned}$$

Therefore,

$$\begin{aligned} \frac{w'(x) - w(x)}{w'(x)} &= \frac{t \left(\frac{F^*(x)}{F_{sw}(x)} - 1 - \log \frac{F^*(x)}{F_{sw}(x)} \right) + \mathcal{O}(t^2)}{1 + t \left(\frac{F^*(x)}{F_{sw}(x)} - 1 \right)} \\ &= \left[t \left(\frac{F^*(x)}{F_{sw}(x)} - 1 - \log \frac{F^*(x)}{F_{sw}(x)} \right) + \mathcal{O}(t^2) \right] \cdot \left[1 - t \left(\frac{F^*(x)}{F_{sw}(x)} - 1 \right) + \mathcal{O}(t^2) \right] \\ &\quad \text{(Claim 3)} \\ &= t \left(\frac{F^*(x)}{F_{sw}(x)} - 1 - \log \frac{F^*(x)}{F_{sw}(x)} \right) + \mathcal{O}(t^2). \end{aligned}$$

In other words, the second term in (28) is

$$\begin{aligned} &\int_{\mathcal{X}} (F_w(x) - F_{sw}(x)) \cdot \phi(x) dx \\ &= - \int_{\mathcal{X}} (F_w(x) - F_{sw}(x)) \cdot \left(\frac{w'(x) - w(x)}{w'(x)} + \mathcal{O}(t^2) \right) dx \quad \text{(Equation (26))} \\ &= \int_{\mathcal{X}} (F_{sw}(x) - F_w(x)) \cdot \left(\frac{w'(x) - w(x)}{w'(x)} \right) dx + \mathcal{O}(t^2) \\ &= t \int_{\mathcal{X}} (F_{sw}(x) - F_w(x)) \left(\frac{F^*(x)}{F_{sw}(x)} - 1 - \log \frac{F^*(x)}{F_{sw}(x)} \right) dx + \mathcal{O}(t^2) \\ &= t \cdot \text{D}_{\text{WIS}}(F^* \| F_{sw}) + \mathcal{O}(t^2), \end{aligned}$$

where the weight function of the weighted Itakura–Saito Divergence is $w = F_{sw} - F_w$. Therefore, as $t \rightarrow 0^+$, the second term in (28) is of the order $\mathcal{O}(t)$. If $\text{D}_{\text{WIS}}(F^* \| F_{sw}) \leq 0$, the second term in (28) will be non-positive (otherwise, it will be positive).

The third term. Taking the limit as $t \rightarrow 0^+$ in (28):

$$\underbrace{\text{d}_{\text{KL}}(F_{sw} \| w')}_{\mathcal{O}(t^2)} + \underbrace{\int_{\mathcal{X}} (F_w(x) - F_{sw}(x)) \cdot \phi(x) dx}_{\mathcal{O}(t)} + t \underbrace{\int_{\mathcal{X}} (F_w(x) - F_{sw}(x)) \log \frac{F_{sw}(x)}{F^*(x)} dx}_{\mathcal{O}(t)} \geq 0.$$

If the middle term is non-positive, the last term should be non-negative:

$$\int_{\mathcal{X}} (F_w(x) - F_{sw}(x)) \log \frac{F_{sw}(x)}{F^*(x)} dx \geq 0. \quad (29)$$

Substituting F^* for g in (21), and using (29), we obtain the desired result

$$d_{\mathcal{P}}(F_{sw}, F^*) \leq d_{\mathcal{P}}(F_w, F^*) - d_{\mathcal{P}}(F_w, F_{sw}).$$

Else, if the middle term is positive (i.e., $\text{D}_{\text{WIS}}(F^* \| F_{sw}) \leq 0$ is not satisfied), the last term may also be non-negative, which means that $d_{\mathcal{P}}(F_{sw}, F^*) \leq d_{\mathcal{P}}(F_w, F^*) - d_{\mathcal{P}}(F_w, F_{sw})$ may also hold. $\square$

The following tools used in the above proof can be proved by Taylor expansion.

Claim 1. For $t, x \in \mathbb{R}_+$, there holds

$$x^t = 1 + t \log x + \frac{1}{2} t^2 (\log x)^2 x^\xi,$$

where $\xi \in (0, t)$.

Claim 2. For $x \in (0, 1)$, there holds:

$$\log(1 - x) = -x - \frac{1}{2} \zeta^2,$$

where $\zeta \in (0, x)$.

Claim 3. For $x \in (-1, 1)$, there holds:

$$\frac{1}{1+x} = 1 - x + \epsilon^2,$$

where ϵ is between 0 and x .

B.6 PROOF OF THEOREM 4

Proof sketch of Charikar et al. (2024). They apply the proof technique from their Theorem 1 to different variables and obtain several inequalities. Subsequently, leveraging the triangle inequality for the ℓ_2 -norm and a uniform convergence argument, they establish the desired result.

Proof sketch of ours. Extending the proof framework from Theorem 2 in (Charikar et al., 2024) is also non-trivial. Specifically, the absence of a triangle inequality for KL divergence necessitates an alternative approach. To address this, we decompose the relevant terms in a manner analogous to the triangle inequality and exhaustively demonstrate that each of the three resulting remainder terms asymptotically converges to zero.

Notations. For a clear presentation, let

$$\begin{aligned} A &= d_{\mathcal{P}}(F_s, F_{sw}) \\ B &= d_{\mathcal{P}}(F_{sw}, F_w) \\ C &= d_{\mathcal{P}}(F_s, F_w) \\ D &= d_{\mathcal{P}}(F^*, F_s) = \varepsilon \\ E &= d_{\mathcal{P}}(F^*, F_{sw}) \\ F &= d_{\mathcal{P}}(F^*, F_w) \\ G &= d_{\mathcal{P}}(F^*, \hat{F}_{sw}) \\ H &= d_{\mathcal{P}}(\hat{F}_{sw}, F_{sw}) \\ I &= d_{\mathcal{P}}(\hat{F}_{sw}, F_w). \end{aligned}$$

Now we start the proof of Theorem 4. A uniform convergence result and two claims used in the proof (Lemma 4, Claim 4 and Claim 5) are provided at the end of the proof.

Proof. Non-realizable weak-to-strong generalization where $F^* \notin V_s$, and we use a finite sample to perform weak-to-strong supervision. Note that by virtue of the range of f^* , f_w and all functions in $\mathcal{F}$ being absolutely bounded, and $d_{\mathcal{P}}$ is also bounded.

Due to $F^* \notin V_s$, we replace F^* with F_s in the final step of proof of Theorem 3, we obtain

$$C \geq A + B. \quad (30)$$

Notice that

$$E = A + D - \underbrace{\int_{\mathcal{X}} (F^*(x) - F_s(x)) \log \frac{F_{sw}(x)}{F_s(x)} dx}_{t_1}, \quad (31)$$

$$F = C + D - \underbrace{\int_{\mathcal{X}} (F^*(x) - F_s(x)) \log \frac{F_w(x)}{F_s(x)} dx}_{t_2}, \quad (32)$$

$$G = E - H - \underbrace{\int_{\mathcal{X}} (\hat{F}_{sw}(x) - F^*(x)) \log \frac{F_{sw}(x)}{\hat{F}_{sw}(x)} dx}_{t_3}. \quad (33)$$

Combining (30) and (31), we get

$$E \leq C + D - B - t_1. \quad (34)$$

By a uniform convergence argument (Lemma 5), we have that with probability at least $1 - \delta$ over the draw of $\{(x_1, y_1), \dots, (x_n, y_n)\}$ that were used to construct $\hat{F}_{sw}$,

$$I \leq B + \underbrace{\mathcal{O}\left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}_s}}{n}}\right)}_{t_4} + \underbrace{\mathcal{O}\left(\sqrt{\frac{\log(1/\delta)}{n}}\right)}_{t_5}. \quad (35)$$

Combining (34) with (35) and we have

$$E \leq C + D - I - t_1 + t_4 + t_5. \quad (36)$$

Combining (32) with (36) and we have

$$E \leq F - I - t_1 + t_2 + t_4 + t_5. \quad (37)$$

Combining (33) with (37) and we have

$$G \leq F - I - H - t_1 + t_2 - t_3 + t_4 + t_5. \quad (38)$$

We replace F^* with $\hat{F}_{sw}$ in the final step of proof of Corollary 1 and obtain:

$$I \geq H + B. \quad (39)$$

Combining (39) with (35) and we have

$$0 \leq H \leq t_4 + t_5 = \mathcal{O}\left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}_s}}{n}}\right) + \mathcal{O}\left(\sqrt{\frac{\log(1/\delta)}{n}}\right). \quad (40)$$

Combining (40) with (38) and we have

$$G \leq F - I - t_1 + t_2 - t_3 + t_4 + t_5. \quad (41)$$

While t_4 and t_5 are known in (35), we analyze t_1 , t_2 and t_3 one by one.

Deal with t_1 . We know that

$$t_1 = \int_{\mathcal{X}} (F^*(x) - F_s(x)) \log \frac{F_{sw}(x)}{F_s(x)} dx.$$

Using Pinsker's inequality and the fact that $\frac{F_{sw}(x)}{F_s(x)} \leq \frac{1}{\gamma}$, we have

$$|t_1| \leq \frac{1}{\gamma} \int_{\mathcal{X}} |F_s(x) - F^*(x)| dx \leq \frac{1}{\gamma} \sqrt{\frac{1}{2} \text{d}_{\text{KL}}(F_s \| F^*)} = \frac{1}{\gamma} \sqrt{\frac{1}{2}} \varepsilon. \quad (42)$$

Therefore,

$$|t_1| = \mathcal{O}(\sqrt{\varepsilon}). \quad (43)$$

Deal with t_2 . The proof for t_2 is similar for t_1 . In particular, replacing F_{sw} with F_w in the above and we can get

$$|t_2| = O(\sqrt{\varepsilon}). \quad (44)$$

Deal with t_3 . We know that

$$t_3 = \int_{\mathcal{X}} (\hat{F}_{sw}(x) - F^*(x)) \log \frac{F_{sw}(x)}{\hat{F}_{sw}(x)} dx.$$

According to Lemma 5, with probability at least $1 - \delta$ over the draw of $(x_1, y_1), \dots, (x_n, y_n)$, we have

$$\left| d_{\mathcal{P}}(\hat{F}_{sw}, F_w) - d_{\mathcal{P}}(F_{sw}, F_w) \right| \leq \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right). \quad (45)$$

According to Claim 4 and (45), we have

$$\left| \int_{\mathcal{X}} \log \frac{F_{sw}(x)}{\hat{F}_{sw}(x)} dx \right| \leq \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right).$$

Since $|\hat{F}_{sw}(x) - F^*(x)|$ is upper bounded, there holds

$$|t_3| = \left| \int_{\mathcal{X}} (\hat{F}_{sw}(x) - F^*(x)) \log \frac{F_{sw}(x)}{\hat{F}_{sw}(x)} dx \right| \leq \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right). \quad (46)$$

Therefore, combining (43), (44) and (46), we have

$$|t_1| + |t_2| + |t_3| \leq O(\sqrt{\varepsilon}) + \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right). \quad (47)$$

Finally, combining (35) and (41) with (40) and (47), we get the result:

$$d_{\mathcal{P}}(F^*, \hat{F}_{sw}) \leq d_{\mathcal{P}}(F^*, F_w) - d_{\mathcal{P}}(\hat{F}_{sw}, F_w) + O(\sqrt{\varepsilon}) + \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right),$$

where in the last inequality, we instantiate asymptotics with respect to $\varepsilon \rightarrow 0$ and $n \rightarrow \infty$. $\square$

Here are some tools used in the above proof.

Lemma 4 (Uniform convergence (forward KL loss)). *Let $(x_1, y_1), \dots, (x_n, y_n)$ be an i.i.d. training sample, where each $x_i \sim \mathcal{P}$ and $y_i = F_w(x_i)$ for a target function F_w . For a fixed strong model representation h_s , let*

$$\begin{aligned} f_{sw} &= \operatorname{argmin}_{f \in \mathcal{F}_s} d_{\mathcal{P}}(F_w, f \circ h_s), \\ \hat{f}_{sw} &= \operatorname{argmin}_{f \in \mathcal{F}_s} \hat{d}_{\mathcal{P}}(F_w, f \circ h_s). \end{aligned}$$

Assume that the range of F_w and functions in $\mathcal{F}_s$ is absolutely bounded. Then, with probability at least $1 - \delta$ over the draw of $(x_1, y_1), \dots, (x_n, y_n)$, we have

$$\left| d_{\mathcal{P}}(F_w, \hat{f}_{sw}) - d_{\mathcal{P}}(F_w, f_{sw}) \right| \leq \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}_s}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right),$$

where $\mathcal{C}_{\mathcal{F}_s}$ is a constant capturing the complexity of the function class $\mathcal{F}_s$.

Proof. The proof is strongly motivated by lemma 4 in Charikar et al. (2024).

Note that

$$\begin{aligned} d_{\mathcal{P}}(F_w, \hat{F}_{sw}) - d_{\mathcal{P}}(F_w, F_{sw}) &= \underbrace{d_{\mathcal{P}}(F_w, \hat{F}_{sw}) - \hat{d}_{\mathcal{P}}(F_w, \hat{F}_{sw})}_a + \\ &\quad \underbrace{\hat{d}_{\mathcal{P}}(F_w, \hat{F}_{sw}) - \hat{d}_{\mathcal{P}}(F_w, F_{sw})}_b + \underbrace{\hat{d}_{\mathcal{P}}(F_w, F_{sw}) - d_{\mathcal{P}}(F_w, F_{sw})}_c. \end{aligned} \quad (48)$$

By the definition of $\hat{f}_{sw}$, the second term $b \leq 0$ in (48). Therefore,

$$\left| d_{\mathcal{P}}(F_w, \hat{F}_{sw}) - d_{\mathcal{P}}(F_w, F_{sw}) \right| \leq |a| + |c|. \quad (49)$$

The terms a and c measure the difference between the empirical risk and true population risk, and can be controlled by a standard uniform convergence argument.

Let $S = \{(x_1, y_1), \dots, (x_n, y_n)\}$, where $x_i \sim \mathcal{P}$ and $y_i = F_w(x_i)$. According to statistical learning theory literature (Bartlett & Mendelson, 2002), it first holds that with probability at least $1 - \delta$,

$$\sup_{f \in \mathcal{F}_s} |\hat{d}_{\mathcal{P}}(F_w, f \circ h_s) - d_{\mathcal{P}}(F_w, f \circ h_s)| \leq O(\mathcal{R}_n(l(\mathcal{F}_s))) + \mathcal{O}\left(\sqrt{\frac{\log(1/\delta)}{n}}\right),$$

where $\mathcal{R}_n(l(\mathcal{F}_s))$ is the *Rademacher complexity* of the loss class of $\mathcal{F}_s$:

$$\mathcal{R}_n(l(\mathcal{F}_s)) = \mathbb{E}_S \mathbb{E}_{\varepsilon_i \sim \{-1, 1\}} \sup_{f \in \mathcal{F}_s} \frac{1}{n} \sum_{i=1}^n \varepsilon_i \cdot \ell(f \circ h_s(x_i), y_i).$$

Notice again that the model output space $\mathcal{Y} = \{y \in \mathbb{R} | 0 < \gamma \leq y \leq 1, \gamma > 0\}$. We can then use the assumption that the range of F_w and $\mathcal{F}_s$ is absolutely bounded, which implies that ℓ is both bounded and Lipschitz in both arguments. This allows us to use the contraction principle in Theorem 4.12 from Ledoux & Talagrand (2013) so as to move from the Rademacher complexity of the loss class $l(\mathcal{F}_s)$ to that of $\mathcal{F}_s$ itself, and claim that with probability at least $1 - \delta$,

$$\sup_{f \in \mathcal{F}_s} |\hat{d}_{\mathcal{P}}(F_w, f \circ h_s) - d_{\mathcal{P}}(F_w, f \circ h_s)| \leq O(\mathcal{R}_n(\mathcal{F}_s)) + \mathcal{O}\left(\sqrt{\frac{\log(1/\delta)}{n}}\right) \quad (50)$$

Finally, the Rademacher complexity $\mathcal{R}_n(\mathcal{F}_s)$ can be upper bounded by a quantity known as the *worst-case Gaussian complexity* of $\mathcal{F}_s$; in any case, for a majority of parametric function classes $\mathcal{F}_s$, this quantity scales as $\sqrt{\frac{\mathcal{C}_{\mathcal{F}_s}}{n}}$ (Bartlett & Mendelson, 2002), where $\mathcal{C}_{\mathcal{F}_s}$ is a constant capturing the inherent complexity of $\mathcal{F}_s$. Plugging this into (50) and considering $f = \hat{f}_{sw}$ or $f = f_{sw}$ in this inequality, we have

$$\begin{aligned} \underbrace{|\hat{d}_{\mathcal{P}}(F_w, \hat{F}_{sw}) - d_{\mathcal{P}}(F_w, \hat{F}_{sw})|}_{|a|} &\leq \mathcal{O}\left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}_s}}{n}}\right) + \mathcal{O}\left(\sqrt{\frac{\log(1/\delta)}{n}}\right), \\ \underbrace{|\hat{d}_{\mathcal{P}}(F_w, F_{sw}) - d_{\mathcal{P}}(F_w, F_{sw})|}_{|c|} &\leq \mathcal{O}\left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}_s}}{n}}\right) + \mathcal{O}\left(\sqrt{\frac{\log(1/\delta)}{n}}\right). \end{aligned}$$

Finally, substitute it into Equation (49) and we can obtain the desired bound. $\square$

Lemma 5 (Uniform convergence (reverse KL loss)). *Let $(x_1, y_1), \dots, (x_n, y_n)$ be an i.i.d. training sample, where each $x_i \sim \mathcal{P}$ and $y_i = F_w(x_i)$ for a target function F_w . For a fixed strong model representation h_s , let*

$$\begin{aligned} f_{sw} &= \operatorname{argmin}_{f \in \mathcal{F}_s} d_{\mathcal{P}}(f \circ h_s, F_w), \\ \hat{f}_{sw} &= \operatorname{argmin}_{f \in \mathcal{F}_s} \hat{d}_{\mathcal{P}}(f \circ h_s, F_w). \end{aligned}$$

Assume that the range of F_w and functions in $\mathcal{F}_s$ is absolutely bounded. Then, with probability at least $1 - \delta$ over the draw of $(x_1, y_1), \dots, (x_n, y_n)$, we have

$$\left| d_{\mathcal{P}}(\hat{F}_{sw}, F_w) - d_{\mathcal{P}}(F_{sw}, F_w) \right| \leq \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}_s}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right),$$

where $\mathcal{C}_{\mathcal{F}_s}$ is a constant capturing the complexity of the function class $\mathcal{F}_s$.

Proof. Swap the order of the two elements in $d_{\mathcal{P}}(\cdot, \cdot)$ and $\hat{d}_{\mathcal{P}}(\cdot, \cdot)$ in the proof of Lemma 4 and we can prove the result. $\square$

Claim 4. Let $f(x), g(x) \in [\gamma, 1]$ where $\gamma > 0$. If there exists $\xi > 0$ such that $\int_{\mathcal{X}} |f(x) - g(x)| dx \leq \xi$, then there holds

$$\int_{\mathcal{X}} |\log f(x) - \log g(x)| dx \leq \frac{1}{\gamma} \xi.$$

Proof. Using the property of the function $\phi(x) = \log x$ (as shown in Figure 4): if $x \in (0, 1]$, then the slope of a line with any two points on the function $\phi(x)$ is bounded.

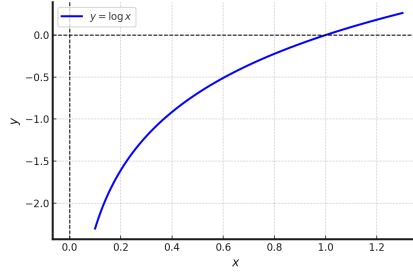


Figure 4: The function $\phi(x) = \log x$.

In particular, we have

$$\begin{aligned} & \int_{\mathcal{X}} |\log f(x) - \log g(x)| dx \\ &= \int_{\mathcal{X}} \left| \frac{\log f(x) - \log g(x)}{f(x) - g(x)} \right| |f(x) - g(x)| dx \\ &\leq \frac{1}{\gamma} \int_{\mathcal{X}} |f(x) - g(x)| dx \\ &\leq \frac{1}{\gamma} \xi. \end{aligned}$$

$\square$

Claim 5. Let $f(x), g(x) \in [\gamma, 1]$ where $\gamma > 0$. If there exists $\xi > 0$ such that $\int_{\mathcal{X}} |\log f(x) - \log g(x)| dx \leq \xi$, then there holds

$$\int_{\mathcal{X}} |f(x) - g(x)| dx \leq \xi.$$

Proof. Using the property of the function $\phi(x) = \log x$ (as shown in Figure 4): if $x \in (0, 1]$, then the slope of a line with any two points on the function $\phi(x)$ is bounded.

In particular, we have

$$\int_{\mathcal{X}} |f(x) - g(x)| dx$$

$$\begin{aligned}
&= \int_{\mathcal{X}} \left| \frac{\log f(x) - \log g(x)}{f(x) - g(x)} \right|^{-1} |\log f(x) - \log g(x)| dx \\
&\leq \int_{\mathcal{X}} |\log f(x) - \log g(x)| dx \\
&\leq \xi.
\end{aligned}$$

□

B.7 EXTENSION OF THEOREM 4

This additional theoretical result follows Corollary 1 in Appendix B.5.

Corollary 2. *Given F^* , F_w and F_{sw} defined above. Let $d_{\mathcal{P}}$ be the output distribution divergence and consider WTSG in Equation (7) using forward KL divergence loss. Consider $\mathcal{F}_s$ that satisfies Assumption 1. If the weighted Itakura–Saito divergence $D_{\text{WIS}}(F^* \| F_{sw}) \leq 0$ with the weight function $w = F_{sw} - F_w$, then we have that with probability at least $1 - \delta$ over the draw of n i.i.d. samples,*

$$d_{\mathcal{P}}(\hat{F}_{sw}, F^*) \leq d_{\mathcal{P}}(F_w, F^*) - d_{\mathcal{P}}(F_w, \hat{F}_{sw}) + \mathcal{O}(\sqrt{\varepsilon}) + \mathcal{O}\left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}_s}}{n}}\right) + \mathcal{O}\left(\sqrt{\frac{\log(1/\delta)}{n}}\right), \quad (51)$$

where $\mathcal{C}_{\mathcal{F}_s}$ is a constant capturing the complexity of the function class $\mathcal{F}_s$, and the asymptotic notation is with respect to $\varepsilon \rightarrow 0, n \rightarrow \infty$.

Proof. For a clear presentation, let

$$\begin{aligned}
A &= d_{\mathcal{P}}(F_{sw}, F_s) \\
B &= d_{\mathcal{P}}(F_w, F_{sw}) \\
C &= d_{\mathcal{P}}(F_w, F_s) \\
D &= d_{\mathcal{P}}(F_s, F^*) = \varepsilon \\
E &= d_{\mathcal{P}}(F_{sw}, F^*) \\
F &= d_{\mathcal{P}}(F_w, F^*) \\
G &= d_{\mathcal{P}}(\hat{F}_{sw}, F^*) \\
H &= d_{\mathcal{P}}(F_{sw}, \hat{F}_{sw}) \\
I &= d_{\mathcal{P}}(F_w, \hat{F}_{sw}).
\end{aligned}$$

The main proof idea follows Appendix B.6. Specifically, we first replace F^* with F_s in the final step of proof of Corollary 1, we obtain

$$C \geq A + B. \quad (52)$$

Notice that

$$E = A + D - \underbrace{\int_{\mathcal{X}} (F_{sw}(x) - F_s(x)) \log \frac{F^*(x)}{F_s(x)} dx}_{t_1}, \quad (53)$$

$$F = C + D - \underbrace{\int_{\mathcal{X}} (F_w(x) - F_s(x)) \log \frac{F^*(x)}{F_s(x)} dx}_{t_2}, \quad (54)$$

$$G = E - H - \underbrace{\int_{\mathcal{X}} (\hat{F}_{sw}(x) - F_{sw}(x)) \log \frac{F^*(x)}{F_{sw}(x)} dx}_{t_3}. \quad (55)$$

Combining (52) and (53), we get

$$E \leq C + D - B - t_1. \quad (56)$$

According to Lemma 4, we have that with probability at least $1 - \delta$ over the draw of $\{(x_1, y_1), \dots, (x_n, y_n)\}$ that were used to construct $\hat{F}_{sw}$,

$$I \leq B + \underbrace{\mathcal{O}\left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}_s}}{n}}\right)}_{t_4} + \underbrace{\mathcal{O}\left(\sqrt{\frac{\log(1/\delta)}{n}}\right)}_{t_5}. \quad (57)$$

Combining (56) with (57) and we have

$$E \leq C + D - I - t_1 + t_4 + t_5. \quad (58)$$

Combining (54) with (58) and we have

$$E \leq F - I - t_1 + t_2 + t_4 + t_5. \quad (59)$$

Combining (55) with (59) and we have

$$G \leq F - I - H - t_1 + t_2 - t_3 + t_4 + t_5. \quad (60)$$

We replace F^* with $\hat{F}_{sw}$ in the final step of proof of Corollary 1 (Recall the fact that $\hat{F}_{sw} \in V_s$ and (29): substituting $\hat{F}_{sw}$ for g in (21), and using (29)), we obtain:

$$I \geq H + B. \quad (61)$$

Combining (61) with (57) and we have

$$0 \leq H \leq t_4 + t_5 = \mathcal{O}\left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}_s}}{n}}\right) + \mathcal{O}\left(\sqrt{\frac{\log(1/\delta)}{n}}\right). \quad (62)$$

Combining (62) with (60) and we have

$$G \leq F - I - t_1 + t_2 - t_3 + t_4 + t_5. \quad (63)$$

While t_4 and t_5 are known in (57), we analyze t_1 , t_2 and t_3 one by one.

Deal with t_1 . We know that

$$t_1 = \int_{\mathcal{X}} (F_{sw}(x) - F_s(x)) \log \frac{F^*(x)}{F_s(x)} dx.$$

Using the fact that $|F_{sw}(x) - F_s(x)| \leq 1$, we have

$$|t_1| \leq \int_{\mathcal{X}} \left| \log \frac{F^*(x)}{F_s(x)} \right| dx = \int_{\mathcal{X}} |\log F_s(x) - \log F^*(x)| dx. \quad (64)$$

According to Pinsker's inequality,

$$\int_{\mathcal{X}} |F_s(x) - F^*(x)| dx \leq \sqrt{\frac{1}{2} \text{d}_{\text{KL}}(F_s \| F^*)} = \sqrt{\frac{1}{2} \varepsilon}. \quad (65)$$

Substitute $f(x) = F_s(x)$, $g(x) = F^*(x)$ and $\xi = \sqrt{\frac{1}{2} \varepsilon}$ into Claim 4 and recall (64), we have

$$|t_1| \leq \frac{1}{\gamma} \sqrt{\frac{1}{2} \varepsilon} = O(\sqrt{\varepsilon}). \quad (66)$$

Deal with t_2 . The proof for t_2 is similar for t_1 . In particular, replacing F_{sw} with F_w in the above and we can get

$$|t_2| = O(\sqrt{\varepsilon}). \quad (67)$$

Deal with t_3 . We know that

$$t_3 = \int_{\mathcal{X}} (\hat{F}_{sw}(x) - F_{sw}(x)) \log \frac{F^*(x)}{\hat{F}_{sw}(x)} dx.$$

According to Lemma 4, with probability at least $1 - \delta$ over the draw of $(x_1, y_1), \dots, (x_n, y_n)$, we have

$$\left| d_{\mathcal{P}}(F_w, \hat{F}_{sw}) - d_{\mathcal{P}}(F_w, F_{sw}) \right| \leq \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right). \quad (68)$$

Notice that

$$\begin{aligned} H &= d_{\mathcal{P}}(F_{sw}, \hat{F}_{sw}) \\ &= d_{\mathcal{P}}(F_w, F_{sw}) - d_{\mathcal{P}}(F_w, \hat{F}_{sw}) + \int_{\mathcal{X}} (F_w(x) + F_{sw}(x)) \log \frac{F_{sw}(x)}{\hat{F}_{sw}(x)} dx. \end{aligned} \quad (69)$$

Substitute (62) and (68) into Equation (69) with the triangle inequality for absolute values, we get

$$\left| \int_{\mathcal{X}} (F_w(x) + F_{sw}(x)) \log \frac{F_{sw}(x)}{\hat{F}_{sw}(x)} dx \right| \leq \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right)$$

Since $|F_w(x) + F_{sw}(x)|$ is bounded, we have

$$\left| \int_{\mathcal{X}} [\log F_{sw}(x) - \log \hat{F}_{sw}(x)] dx \right| \leq \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right).$$

Using Claim 5, we have

$$\left| \int_{\mathcal{X}} (\hat{F}_{sw}(x) - F_{sw}(x)) dx \right| \leq \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right).$$

Since $\left| \log \frac{F^*(x)}{\hat{F}_{sw}(x)} \right|$ is bounded, there holds

$$|t_3| = \left| \int_{\mathcal{X}} (\hat{F}_{sw}(x) - F_{sw}(x)) \log \frac{F^*(x)}{\hat{F}_{sw}(x)} dx \right| \leq \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right). \quad (70)$$

Therefore, combining (66), (67) and (70), we have

$$|t_1| + |t_2| + |t_3| \leq \mathcal{O}(\sqrt{\varepsilon}) + \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right). \quad (71)$$

Finally, combining (57) and (63) with (62) and (71), we get the result:

$$d_{\mathcal{P}}(\hat{F}_{sw}, F^*) \leq d_{\mathcal{P}}(F_w, F^*) - d_{\mathcal{P}}(F_w, \hat{F}_{sw}) + \mathcal{O}(\sqrt{\varepsilon}) + \mathcal{O} \left(\sqrt{\frac{\mathcal{C}_{\mathcal{F}}}{n}} \right) + \mathcal{O} \left(\sqrt{\frac{\log(1/\delta)}{n}} \right),$$

where in the last inequality, we instantiate asymptotics with respect to $\varepsilon \rightarrow 0$ and $n \rightarrow \infty$. $\square$

C FURTHER DETAILS AND RESULTS OF EXPERIMENTS

C.1 TRAINING DETAILS OF EXPERIMENTS IN LANGUAGE MODELS

The dataset is randomly divided into three distinct subsets:

- 4K samples (ground truth): They are used to fine-tune weak and strong base language models;

- 4K samples (held-out set): These labels are predicted by the weak model and used to provide weak supervision for training the strong model;
- 4K samples (the remaining): They are used for testing and evaluating the performance of all models.

When fine-tuning the weak-to-strong models, we follow (Yang et al., 2025) to set the batch size to 32, learning rate to 10^{-5} , `max_seq_len` to 512. The training epoch is set to 1 to avoid overfitting. All experiments are conducted on NVIDIA A100 80G.

C.2 WEAK-TO-STRONG TRAINING PROCEDURE IN SYNTHETIC EXPERIMENTS

We explore two methods to generate the weak and strong representations, which follows the experimental setting in (Charikar et al., 2024).

- **Pre-training.** We begin by randomly sampling T fine-tuning tasks $f_1^*, \dots, f_T^* \in \mathcal{F}_s$. For each $t \in \{1, \dots, T\}$, we generate N_r data $\{x_j^{(t)}\}_{j=1}^{N_r}$ where $x_j^{(t)} \sim \mathcal{P}$. Let the representations $h_w, h_s : \mathbb{R}^8 \rightarrow \mathbb{R}^{16}$ be 2-layer and 8-layer MLP with ReLU activations, respectively. And $h_w \in \mathcal{H}_w$, $h_s \in \mathcal{H}_s$. We obtain h_w and h_s via gradient descent on the representation parameters to find the minimizer of output distribution divergence loss. Specifically, We use Equation (4) as the loss function on T tasks:

$$h_l = \operatorname{argmin}_{h \in \mathcal{H}_l} \frac{1}{T} \sum_{t=1}^T \hat{d}_{\mathcal{P}}(f_t^* \circ h, f_t^* \circ h^*), \quad (72)$$

where $l \in \{w, s\}$, $T = 10$, and $N_r = 2000$. Additionally, the realizable setting (Corollary 1) is considered by explicitly setting $h_s = h^*$, and only obtaining h_w as above.

- **Perturbations.** As an alternative, we directly perturb the parameters of h^* to obtain the weak and strong representations. Specifically, we add independent Gaussian noise $\mathcal{N}(0, \sigma_s^2)$ to every parameter in h^* to generate h_s . Similarly, we perturb h^* with $\mathcal{N}(0, \sigma_w^2)$ to generate h_w . To ensure the strong representation h_s is closer to h^* than h_w , we set $\sigma_s = 0.1$ and $\sigma_w = 9$.

Weak Model Fine-tuning. After obtaining h_w and h_s , we fix these representations and train weak models on new fine-tuning tasks. We randomly sample M new fine-tuning tasks $f_1^*, \dots, f_M^* \in \mathcal{F}_s$, and generate data $\{x_j^{(i)}\}_{j=1}^{N_f}$, where $x_j^{(i)} \sim \mathcal{P}$. For each task $i = \{1, \dots, M\}$, the weak model is trained through:

$$f_w^{(i)} = \operatorname{argmin}_f \frac{1}{M} \sum_{i=1}^M \hat{d}_{\mathcal{P}}(f_t^* \circ h^*, f \circ h_w), \quad (73)$$

where $M = 100$, $N_f = 2000$. Here, the representation parameters h_w are frozen, and $f_w^{(i)}$ is learned via gradient descent. Weak models are thus trained on true data.

Weak-to-Strong Supervision. Using the trained weak models, we generate weakly labeled datasets for each fine-tuning task. Specifically, for each $i \in \{1, \dots, M\}$, we generate $\{\tilde{x}_j^{(i)}\}_{j=1}^{N_f}$ where $\tilde{x}_j^{(i)} \sim \mathcal{P}$. The strong models are then trained on these weakly labeled datasets by solving the following optimization problem using reverse KL divergence loss for each task $i \in \{1, \dots, M\}$:

$$f_{sw}^{(i)} = \operatorname{argmin}_{f \in \mathcal{F}} \hat{d}_{\mathcal{P}}(f \circ h_s, f_w^{(i)} \circ h_w). \quad (74)$$

At this stage, the weak-to-strong training procedure is complete.