
Blessing of Class Diversity in Pre-training

Yulai Zhao
Princeton University

Jianshu Chen
Tencent AI Lab

Simon S. Du
University of Washington

Abstract

This paper presents a new statistical analysis aiming to explain the recent superior achievements of the pre-training techniques in natural language processing (NLP). We prove that when the classes of the pre-training task (e.g., different words in the masked language model task) are sufficiently diverse, in the sense that the least singular value of the last linear layer in pre-training (denoted as $\tilde{\nu}$) is large, then pre-training can significantly improve the sample efficiency of downstream tasks. Specially, we show the transfer learning excess risk enjoys an $O\left(\frac{1}{\tilde{\nu}\sqrt{n}}\right)$ rate, in contrast to the $O\left(\frac{1}{\sqrt{m}}\right)$ rate in the standard supervised learning. Here, n is the number of pre-training data and m is the number of data in the downstream task, and typically $n \gg m$. Our proof relies on a vector-form Rademacher complexity chain rule for disassembling composite function classes and a modified self-concordance condition. These techniques can be of independent interest.

1 INTRODUCTION

Pre-training refers to training a model on a few or many tasks to help it learn parameters that can be used in other tasks. For example, in natural language processing (NLP), one first pre-trains a complex neural network model to predict masked words (masked language modeling), and then fine-tunes the model on downstream tasks, e.g., sentiment analysis (Devlin et al., 2019).

Recently, the pre-training technique has revolutionized the NLP area. Models based on this technique have dramatically improved the performance for numerous downstream tasks (Devlin et al., 2019; Radford et al., 2018; Yang et al.,

2019; Clark et al., 2020; Lan et al., 2020; Liu et al., 2020).

Despite the large body of empirical work on pre-training, satisfactory theories are still lacking, especially theories that can explain the success of pre-training in NLP. Existing theories often rely on strong distributional assumptions (Lee et al., 2021), smoothness conditions (Robinson et al., 2020) or noise-robustness conditions (Bansal et al., 2021) to relate the pre-training task(s) to downstream tasks. These assumptions are often hard to verify.

A line of work studied *multi-task* pre-training (Caruana, 1997; Baxter, 2000; Maurer et al., 2016; Du et al., 2021; Tripuraneni et al., 2021, 2020; Thekumparampil et al., 2021). In particular, recently, researchers have identified a new condition, the *diversity* of pre-training tasks, which has been shown to be crucial to allowing pre-trained models to be useful for downstream tasks. See Section 2 for more detailed discussions on related work.

Unfortunately, this line of theory cannot be used to explain the success of pre-training in NLP. The theory of multi-task pre-training requires *a large number of diverse tasks*, e.g., the number of tasks needs to be larger than the last layer’s input dimension (a.k.a. embedding dimension), which is typically 768, 1024, or 2048 (Devlin et al., 2019). However, in NLP pre-training, there are only a few, if not one, pre-training tasks. Therefore, we need a new theory that applies to this setting.

Since in NLP pre-training, we do not have multiple tasks, we propose to study *the blessing of multiple classes*. Concretely, consider the Masked Language Model (MLM) pre-training task in NLP. In such a pre-training task, we have a large collection of sentences (e.g. from Wikipedia). During the pre-training phase, we randomly mask a few words in each sentence and predict the masked words using the remaining words in this sentence. This pre-training task is a *multi-class classification problem* where the number of classes is about 30K when using byte-pair-encoding (BPE) sub-word units.¹ Note that this number is much larger than the embedding dimension (768, 1024, or 2048).

¹This is a standard setting in the BERT model (Devlin et al., 2019) and is widely adopted as a common practice. By breaking down the English words into BPE sub-word units, it could drastically increase the coverage of the English language by using a relatively small (32768) vocabulary.

In this paper, we develop a new statistical analysis aiming to explain the success of pre-training for NLP. The key notion of our theory is the *diversity of classes*, which serves a similar role as the *diversity of tasks* in multi-task pre-training theory (Du et al., 2021; Tripuraneni et al., 2021). We summarize our contributions below.

First, we define a new notion, *diversity of classes*, which is the least singular value of the last linear layer in pre-training. We prove finite-sample bounds to show that for the cross-entropy loss, if the diversity of classes is large, then pre-training on a single task provably improves the statistical efficiency of the downstream tasks. We give concrete bounds on linear representation and deep neural networks to showcase our general theoretical results. To our knowledge, this is the first set of theoretical results that demonstrates the statistical gain of the standard practice of NLP pre-training, without strong distributional or smoothness conditions.

Second, from a technical point of view, previous theoretical work on multi-task learning (Du et al., 2021; Tripuraneni et al., 2020) builds on scalar output and thus could not apply to multi-class tasks (e.g., cross-entropy loss). We introduce a vector-form Rademacher complexity chain rule for disassembling composite function classes based on vector-form Rademacher contraction property (Maurer, 2016). This generalizes the scalar-form chain rule in Tripuraneni et al. (2020). Furthermore, we adopt the *modified self-concordance* condition to show that the least singular value of the last linear layer serves as a diversity parameter for cross-entropy loss. We believe our techniques can be useful in other problems.

Organization. This paper is organized as follows. In Section 2, we review the related work. In Section 3, we formally describe the problem setup and introduce the necessary definitions. In Section 4, we state our main Theorem 4.2 then instantiate it with several settings. We conclude and discuss some interesting future directions in Section 5. All proofs are deferred to Appendix A. In Appendix B, we present some preliminary empirical results on how our theory inspires new regularization techniques.

2 RELATED WORK

Here we mostly focus on the theoretical aspects of pre-training. While there is a long list of work demonstrating the empirical success of self-supervised learning, there are only a few papers that study its theoretical aspects. One line of work studied the theoretical properties of *contrastive learning* (Saunshi et al., 2019; Tosh et al., 2021), which is a different setting considered in this paper. The most relevant one is by Lee et al. (2021) which showed that if the input data and pre-training labels were independent (conditional on the downstream labels), then pre-training provably im-

proved statistical efficiency. However, this conditional independence assumption rarely holds in practice. For example, in the question-answering task, this assumption implies that given the answer, the question sentence and the masked word are independent. Robinson et al. (2020) assumed the Central Condition and a smoothness condition that relates the pretraining task and the downstream task. Bansal et al. (2021) related generalization error of self-supervised learning to the noise-stability and rationality. However, it is difficult to verify the assumptions in these papers.

A recent line of theoretical work studied multi-task pre-training (Du et al., 2021; Tripuraneni et al., 2021, 2020; Thekumparampil et al., 2021) in which a notion, diversity, has been identified to be the key that enables pre-training to improve statistical efficiency. Experiments also supported the idea that increasing the diversity of the training data helps generalization (Zhang et al., 2022).

Theories on multi-task pre-training generally require a large number of diverse tasks, and thus are not applicable to NLP, as we have mentioned. In comparison, we study single-task multi-class pre-training which is different from theirs. Du et al. (2021) noted that their results allowed an easy adaptation to multi-class settings (see Remark 6.2 therein). However, they only focused on quadratic loss with one-hot labels for multi-class classification. Instead, we study the commonly used cross-entropy loss.

While their analyses do not imply results in our setting, our theoretical analyses are inspired by this line of work.

3 PRELIMINARIES

In this section, we introduce the necessary notations, the problem setup, and several model-dependent quantities used in pre-training and downstream task learning.

3.1 Notations and Setup

Notations Let $[n] = \{1, 2, \dots, n\}$. We use $\|\cdot\|$ or $\|\cdot\|_2$ to denote the ℓ_2 norm of a vector. Let $\mathcal{N}(\mu, \sigma_2)$ be the one-dimensional Gaussian distribution. For a matrix $\mathbf{W} \in \mathbb{R}^{m \times n}$, let $\|\mathbf{W}\|_{1,\infty} = \max_q (\sum_p |\mathbf{W}_{q,p}|)$ and $\|\mathbf{W}\|_{\infty \rightarrow 2}$ be the induced ∞ -to-2 operator norm. We use the standard $O(\cdot)$, $\Omega(\cdot)$ and $\Theta(\cdot)$ notation to hide universal constant factors, and use $\tilde{O}(\cdot)$ to hide logarithmic factors. We also use $a \lesssim b$ to indicate $a = O(b)$.

Problem setup This work is in line with previous transfer learning theories (Du et al., 2021; Tripuraneni et al., 2020) that first pre-train on a large corpus to get a good representation, which, could be future utilized by various downstream tasks. Formally, the procedure is divided into two stages: the pre-training stage to find a representation function and the downstream training stage to obtain a predictor for the downstream task. In both stages, we use $\hat{R}$

to represent empirical risk and use R to represent expected loss.

In the first stage, we have one pre-training task with n samples, $\{x_i^p, y_i^p\}_{i=1}^n$, where $x_i^p \in \mathcal{X}^p \subset \mathbb{R}^d$ is the input and $y_i^p \in \{0, 1\}^{k-1}$ is the one-hot label for k -class classification (if y_i^p is all-zero then it represents the k -th class).² For instance, in masked language modeling, the input of each sample is a sentence with one word masked out, and the label is the masked word.³ k in this example is the size of the vocabulary ($\approx 30K$). We aim to obtain a good representation function $\hat{h}$ within a function class $\mathcal{H} \subset \{\mathbb{R}^d \rightarrow \mathbb{R}^r\}$ where r is the embedding dimension (often equals to 768, 1024, 2048 in NLP pre-training). For example, one popular choice of the representation function $\hat{h}$ in NLP applications is the Transformer model and its variants (Vaswani et al., 2017; Devlin et al., 2019). On top of the representation, we predict the labels using function f^p within function class $\mathcal{F}^p \subset \{\mathbb{R}^r \rightarrow \mathbb{R}^{k-1}\}$.

To train the representation function and predictor in pre-training stage, we consider the Empirical Risk Minimization (ERM) procedure

$$\begin{aligned} \hat{h} &= \arg \min_{h \in \mathcal{H}} \min_{f^p \in \mathcal{F}^p} \hat{R}_p(f^p, h) \\ &\triangleq \arg \min_{h \in \mathcal{H}} \min_{f^p \in \mathcal{F}^p} \frac{1}{n} \sum_{i=1}^n \ell(f^p \circ h(x_i^p), y_i^p) \end{aligned}$$

where ℓ is the loss function. We overload the notation for both the pre-training task and the downstream task, i.e., for pre-training, $\ell : \mathbb{R}^{k-1} \times \{0, 1\}^{k-1} \rightarrow \mathbb{R}$ and for the downstream task, $\ell : \mathbb{R}^{k'-1} \times \{0, 1\}^{k'-1} \rightarrow \mathbb{R}$. e.g., cross-entropy: $\ell(\hat{y}; y) = -y^\top \hat{y} + \log(1 + \sum_{s=1}^{k-1} \exp(\hat{y}_s))$.

Now for the downstream task, we assume there are m samples $\{x_i^d, y_i^d\}_{i=1}^m$. Note that $x_i^d \in \mathcal{X}^d \subset \mathbb{R}^d$ is the input and $y_i^d \in \{0, 1\}^{k'-1}$ is the one-hot label for k' -class classification.⁴ Note that in most real-world applications, we have $n \gg m$ and $k \gg k'$. For example, in sentiment analysis, $k' = 2$ ("positive" or "negative"). A widely studied task SST-2 (Wang et al., 2019) has $m \approx 67K$, which is also generally much smaller than the pre-training corpus (e.g., $n > 100M$ samples).

For the downstream task, we fix the representation function learned from the pre-training task and train the task-dependent predictor within $\mathcal{F}^d \subset \{\mathbb{R}^r \rightarrow \mathbb{R}^{k'-1}\}$:

$$\hat{f}^d = \arg \min_{f^d \in \mathcal{F}^d} \hat{R}_d(f^d, \hat{h})$$

²We assume only one pre-training task for the ease of presentation. It is straightforward to generalize our results to multiple pre-training tasks.

³Here we say only one word being masked only for the ease of presentation. It is straightforward to generalize our results to the case where multiple words are masked out.

⁴For simplicity, we assume we only have one downstream task. Our theoretical results still apply if we have multiple downstream tasks.

$$\triangleq \arg \min_{f^d \in \mathcal{F}^d} \frac{1}{m} \sum_{i=1}^m \ell(f^d \circ \hat{h}(x_i^d), y_i^d).$$

Therefore, our predictor for the downstream task consists of a pair $(\hat{f}^d, \hat{h})$. We use the following risk to measure the performance of predictor and representation

Transfer Learning Risk $\triangleq$

$$R_d(\hat{f}^d, \hat{h}) = \mathbb{E}_{x^d, y^d} [\ell(g^d(x^d), y^d)]$$

where we define

$$R_d(\hat{f}^d, \hat{h}) \triangleq \mathbb{E}_{x^d, y^d} [\ell(\hat{f}^d \circ \hat{h}(x^d), y^d)]$$

as the expected loss (the expectation is over the distribution of the downstream task), and

$$g^d = \arg \min_{g \in \{\mathbb{R}^d \rightarrow \mathbb{R}^{k'-1}\}} \mathbb{E}_{x^d, y^d} [\ell(g(x^d), y^d)]$$

is the optimal predictor for the downstream task.

In our analysis, we also need to use the following term to characterize the quality of pre-training

$$\text{Pre-training Risk} \triangleq R_p(\hat{f}^p, \hat{h}) = \mathbb{E}_{x^p, y^p} [\ell(g^p(x^p), y^p)],$$

where

$$R_p(\hat{f}^p, \hat{h}) \triangleq \mathbb{E}_{x^p, y^p} [\ell(\hat{f}^p \circ \hat{h}(x^p), y^p)]$$

is the expected loss, and

$$g^p = \arg \min_{g \in \{\mathbb{R}^d \rightarrow \mathbb{R}^{k-1}\}} \mathbb{E}_{x^p, y^p} [\ell(g(x^p), y^p)]$$

is the optimal predictor for the pre-training task.

Following the existing work on representation learning (Maurer et al., 2016; Du et al., 2021; Tripuraneni et al., 2020), throughout the paper, we make the following realizability assumption, which is also a standard assumption in the classical PAC learning framework (Shalev-Shwartz and Ben-David, 2014).

Assumption 3.1 (Realizability). There exist $h \in \mathcal{H}$, $f^p \in \mathcal{F}^p$, $f^d \in \mathcal{F}^d$ such that $g^p = f^p \circ h$ and $g^d = f^d \circ h$.

This assumption posits that the representation class and the task-dependent prediction classes are sufficiently expressive to contain the optimal functions. Importantly, the pre-training and downstream tasks share a *common* optimal representation function h . This assumption formalizes the intuition that pre-training learns a good representation that is also useful for downstream tasks.

As for the setting that is of most interest to NLP pre-training, where the loss function ℓ is cross-entropy, $\mathcal{F}^p$ and $\mathcal{F}^d$ are sets of linear functions, we make the following assumption on both pre-training and downstream tasks to describe how the underlying data are generated.

Assumption 3.2 (Multinomial Logistic Data). For a K -class classification task with q samples, $\{x_i, y_i\}_{i=1}^q$, where $x_i \in \mathcal{X}$ is the input and $y_i \in \{0, 1\}^{K-1}$ is the one-hot label. Let f and h be the true underlying predictor layer and representation function. Then the output is $f \circ h(x) \in \mathbb{R}^{K-1}$. Assume each label $\{y\}_i$ is generated from a conditional distribution of a multinomial logistic regression model: $y \sim \mathcal{P}(\cdot | f \circ h(x))$,

$$\mathcal{P}(y | f \circ h(x)) = e^{y^\top f \circ h(x) - \Phi(f \circ h(x))}$$

where $\Phi(x) = \log(1 + \sum_{s=1}^{K-1} e^{x_s})$, $x \in \mathbb{R}^{K-1}$ and y is a one-hot label.

Remark 3.3. It is straight forward to see that $\mathcal{P}(y | f \circ h(x))$ is normalized to 1.

Intuitively, the assumption states that the data used for classification follow a multinomial logistic regression structure.

3.2 Task-Relatedness and Diversity

We shall use the following definitions, which are natural analogies of those in [Tripuraneni et al. \(2020\)](#) for multi-task transfer learning. Being in the same framework of developing the diversity of the pre-training phase, [Tripuraneni et al. \(2020\)](#) aimed at improving correlations between K separate and easy tasks, while we show the diversity across various classes in a single but comprehensive pre-training task has prominent effects.

To measure the ‘‘closeness’’ between the learned representation and true underlying feature representation, we use the following metric, following [Tripuraneni et al. \(2020\)](#)

Definition 3.4. Let $h \in \mathcal{H}$ be the optimal representation function and $h' \in \mathcal{H}$ be any representation function. Let $f^p \in \mathcal{F}^p$ be the optimal pre-training predictor on top of h . The pre-training representation difference is defined as:

$$d_{\mathcal{F}^p, f^p}(h'; h) = \inf_{f' \in \mathcal{F}^p} \mathbb{E}_{x^p, y^p} [\ell(f' \circ h'(x^p), y^p) - \ell(f^p \circ h(x^p), y^p)]$$

where the expectation is over the pre-training data distribution.

Intuitively, this measures the performance difference between the optimal predictor and the best possible predictor given a representation h' .

For transfer learning, we also need to introduce a similar concept on the downstream task.

Definition 3.5. Let $h \in \mathcal{H}$ be the optimal representation function and $h' \in \mathcal{H}$ be any representation function. For the downstream task, for a function class $\mathcal{F}^d$, let $f^d \in \mathcal{F}^d$ be the optimal pre-training predictor on top of a specific h .

We define the worst-case representation difference between h and $h' \in \mathcal{H}$ as:

$$d_{\mathcal{F}^d}(h'; h) = \sup_{f^d \in \mathcal{F}^d} \inf_{f' \in \mathcal{F}^d} \mathbb{E}_{x^d, y^d} [\ell(f' \circ h'(x^d), y^d) - \ell(f^d \circ h(x^d), y^d)]$$

where the expectation is over the data distribution of the downstream task. Here, the supremum is taken over $\{f^d | f^d \in \mathcal{F}^d, f^d \text{ is the optimal predictor on } h \in \mathcal{H}\}$.

We now introduce the key notion of *diversity*, which measures how well a learned representation, say h' , from the pre-training task can be transferred to the downstream task.

Definition 3.6. Let $h \in \mathcal{H}$ be the optimal representation function. Let $f^p \in \mathcal{F}^p$ be the optimal pre-training predictor on top of h . The **diversity parameter** $\nu > 0$ is the largest constant that satisfies

$$d_{\mathcal{F}^d}(h'; h) \leq \frac{d_{\mathcal{F}^p, f^p}(h'; h)}{\nu}, \forall h' \in \mathcal{H}. \quad (1)$$

The interpretation of ν is that it serves as a task-relatedness parameter. While Definition 3.4- 3.6 are naturally defined from inspecting the pre-training procedure, it is not trivial to use these definitions to derive statistical guarantees. In particular, one of our key technical challenge is to show the least singular value of the last linear layer serves as a lower bound of the diversity parameter when $\mathcal{F}^p$ and $\mathcal{F}^d$ are linear function classes.

3.3 Model Complexities

Lastly, we need to introduce some notions to measure the complexity of the function classes considered. In this paper, we consider Gaussian complexity which quantifies the extent to which the function in the class $\mathcal{Q}$ can be correlated with a noise sequence of length $n \times r$.

Definition 3.7 (Gaussian Complexity). Let μ be a probability distribution on a set $\mathcal{X} \subset \mathbb{R}^d$ and suppose that $x_1, \dots, x_n$ are independent samples selected according to μ . Let $\mathcal{Q}$ be a class of functions mapping from $\mathcal{X}$ to $\mathbb{R}^r$. Define random variable

$$\hat{G}_n(\mathcal{Q}) = \mathbb{E}_{g_{ki} \sim \mathcal{N}(0,1)} \left[\sup_{q \in \mathcal{Q}} \frac{1}{n} \sum_{k=1}^r \sum_{i=1}^n g_{ki} q_k(x_i) \right] \quad (2)$$

as the empirical Gaussian complexity, where $q_k(x_i)$ is the k -th coordinate of the vector-valued function $q(x_i)$, g_{ki} ($k \in [r], i \in [n]$) are independent standard normal random variables. The Gaussian complexity of $\mathcal{Q}$ is $G_n(\mathcal{Q}) = E_\mu \hat{G}_n(\mathcal{Q})$.

Our main results are stated in terms of the Gaussian complexity. In Section 4.4 and 4.5 we will plug in existing results of the Gaussian complexity of certain function classes to obtain concrete bounds.

We will need the following worst-case Gaussian complexity for the pre-training predictor within $\mathcal{F}^p$

$$\bar{G}_n(\mathcal{F}^p) = \max_{h(x_1), \dots, h(x_n)} \hat{G}_n(\mathcal{F}^p | h(x^p)), \quad (3)$$

here $h \in \mathcal{H}$ and $x^p = x_1, \dots, x_n \in \mathcal{X}^p$. Similarly we define $\bar{G}_m(\mathcal{F}^d)$ as the worst-case Gaussian complexity for the downstream predictor within $\mathcal{F}^d$.

We note that a closely related notion is Rademacher complexity. The empirical Rademacher complexity and Gaussian complexity only differ by a log factor (Ledoux and Talagrand, 1991). We use Gaussian complexity in this work for its benign properties brought by Gaussian distribution.

4 MAIN RESULTS

In this section, we present our main theoretical results. In Section 4.1 we present an analysis in terms of the diversity parameter for general loss function under certain regularity conditions. In Section 4.2, we specialize the result to a setting that is most relevant to NLP pre-training applications, where $\mathcal{F}^p$ and $\mathcal{F}^d$ are sets of linear functions and the loss is cross-entropy. In this particular case, our key result will show that one can use the singular value of the last layer to bound the diversity parameter. In Section 4.4 and 4.5 we instantiate our bounds on two concrete representation function classes: linear subspace and multi-layer network to showcase our main results.

4.1 Main Theorem

In this subsection, we present our generic end-to-end transfer learning guarantee for multi-class transfer learning problems. We do not impose any specific function class formulations. Throughout this subsection, we only make the following mild regularity assumptions to make our results general.

Assumption 4.1 (Regularity Conditions). We assume the following regularity conditions hold:

- In pre-training, $\ell(\cdot, \cdot)$ is B^p -bounded, and $\ell(\cdot, y)$ is L^p -Lipschitz for all y .
- In downstream task, $\ell(\cdot, y)$ is B^d -bounded and L^d -Lipschitz for all y .
- Any predictor $f \in \mathcal{F}^p$ is $L(\mathcal{F}^p)$ -Lipschitz with respect to the Euclidean distance.
- Predictors are bounded: $\|f \circ h(x)\| \leq D_{\mathcal{X}^p}$ for any $x \in \mathcal{X}^p, h \in \mathcal{H}, f \in \mathcal{F}^p$. Similarly $\|f \circ h(x)\| \leq D_{\mathcal{X}^d}$ for any $x \in \mathcal{X}^d, h \in \mathcal{H}, f \in \mathcal{F}^d$.

Specifically, one can show that common task-dependent losses satisfy these conditions. For example, when ℓ

is the cross-entropy loss for k -class classification (cf. Section 4.2), we prove that ℓ is $\sqrt{k-1}$ -Lipschitz and $D_{\mathcal{X}}$ -bounded where $\mathcal{X}$ denotes the input data domain.

Under these assumptions, we have the following quantitative guarantee.

Theorem 4.2. *Under Assumption 3.1 and 4.1, for a given fixed failure probability δ , with probability at least $1 - \delta$ we have the Transfer Learning Risk upper bounded by:*

$$O\left(\frac{1}{\nu} \left\{ L^p \left[\log(n) [L(\mathcal{F}^p) G_n(\mathcal{H}) + \bar{G}_n(\mathcal{F}^p)] + \frac{\sqrt{k} D_{\mathcal{X}^p}}{n^2} \right] + B^p \sqrt{\frac{\log(1/\delta)}{n}} \right\} + L^d \bar{G}_m(\mathcal{F}^d) + B^d \sqrt{\frac{\log(1/\delta)}{m}} \right).$$

The first line comes from the pre-training ERM procedure and it accounts for the error of using an approximate optimal representation $\hat{h} \approx h$. The second line characterizes the statistical error of learning the downstream-task predictor f^d from m samples. Note the diversity parameter appears in the denominator, which relates the pre-training risk to the transfer learning risk. Theorem 4.2 shows the risk would be small if the Gaussian complexities are small. We expect that $G_n(\mathcal{H}) \gg \bar{G}_m(\mathcal{F}^d)$ since $\mathcal{H}$ is often expressive representation functions, while $\mathcal{F}^d$ is linear classifiers generally. We will show concrete examples where $G_n(\mathcal{H})$ and $\bar{G}_m(\mathcal{F}^p)$ are $O(\sqrt{1/n})$ and $\bar{G}_m$ scales as $O(\sqrt{1/m})$. We believe this theorem applies broadly beyond the concrete settings considered in this paper.

In comparison with previous results, transfer learning risk analyses in (Du et al., 2021; Tripuraneni et al., 2020) focus on scalar output. Their results cannot be applied to multi-class transfer learning tasks. In Theorem 4.2, we generalize the analysis in (Tripuraneni et al., 2020) to handle multi-class classification where the output is high dimensional (number of classes). Technically, in the proof, we introduce a vector-form Rademacher complexity chain rule for disassembling composite function classes by making use of the vector-form Rademacher contraction property (Maurer, 2016).

4.2 Multi-class Classification with Cross-entropy Loss

Now we specialize the general results to the setting that is of most interest to NLP pre-training, where the loss function ℓ is cross-entropy and the $\mathcal{F}^p$ and $\mathcal{F}^d$ are sets of linear functions. This choice is consistent with the NLP pre-training: e.g., BERT (Devlin et al., 2019) uses transformers as the representation learning function class $\mathcal{H}$ and uses word-embedding matrices as $\mathcal{F}^p$.

Formally we define

$$\mathcal{F}^p = \{f | f(z) = \alpha^\top z, \alpha \in \mathbb{R}^{r \times (k-1)},$$

$$\begin{aligned} & \|\alpha_s\| \leq c_1 \text{ for all } s \in [k-1], \|\alpha^\top z\| \leq c_2\} \\ \mathcal{F}^d = \{f|f(z) = \alpha^\top z, \alpha \in \mathbb{R}^{r \times (k'-1)}, \\ & \|\alpha_s\| \leq c_0 \text{ for all } s \in [k'-1], \|\alpha^\top z\| \leq c_3\} \end{aligned}$$

where c_0, c_1, c_2 and c_3 are some positive constants. Then the regularity conditions are instantiated as:

- Pre-training loss $\ell(\cdot, y)$ is $\sqrt{k-1}$ -Lipschitz and $B^p = D_{\mathcal{X}^p}$ -bounded.
- Downstream loss is $\sqrt{k'-1}$ -Lipschitz and $B^d = D_{\mathcal{X}^d}$ -bounded.
- Any $f \in \mathcal{F}^p$ is $L(\mathcal{F}^p) = c_1\sqrt{k-1}$ -Lipschitz w.r.t. the ℓ_2 distance.

Next, we discuss our main assumption that relates the diversity parameter to a concrete quantity of the last linear layer.

Assumption 4.3 (Lower Bounded Least Eigenvalue). Let the optimal linear predictor at the last layer for pre-training be $\alpha^p \in \mathbb{R}^{r \times (k-1)}$, $\tilde{\nu} \triangleq \sigma_r(\alpha^p (\alpha^p)^\top) > 0$ where σ_r is the r -biggest eigenvalue.

Similar assumptions have been used in multi-task representation learning (Du et al., 2021; Tripuraneni et al., 2021, 2020), and are shown to be necessary (Maurer et al., 2016; Du et al., 2021). Different from their versions, our assumption is tailored for the multi-class classification setting. We provide proof sketches on how $\tilde{\nu}$ serves as a lower bound for the diversity parameter ν (cf. Lemma 4.5) in Section 4.3, where we introduce new techniques for analysis.

Intuitively, this assumption ensures that the pre-training task matrix spans the entire r -dimensional space and thus covers the output of the optimal representation $h(\cdot) \in \mathbb{R}^r$. This is quantitatively captured by the $\sigma_r(\alpha^p (\alpha^p)^\top)$, which measures how spread out these vectors are in $\mathbb{R}^r$.

We now state our theorem for this specific setting.

Theorem 4.4. *Under Assumption 3.1, 3.2, 4.3, with probability at least $1 - \delta$ we have the Transfer Learning Risk upper bounded by:*

$$\begin{aligned} & O\left(\frac{1}{\tilde{\nu}} \left\{ \sqrt{k} \left[\log(n) [\sqrt{k} G_n(\mathcal{H}) + \bar{G}_n(\mathcal{F}^p)] + \frac{\sqrt{k} D_{\mathcal{X}^p}}{n^2} \right] \right. \right. \\ & \left. \left. + D_{\mathcal{X}^p} \sqrt{\frac{\log(1/\delta)}{n}} \right\} + \sqrt{k'} \mathbb{E}_{\mathcal{X}^d} \hat{G}_m(\mathcal{F}^d | \hat{h} \circ x^d) \right. \\ & \left. + \sigma \sqrt{\frac{\log(1/\delta)}{m}} + D_{\mathcal{X}^d} \sqrt{\frac{\log(1/\delta)}{m}} \right) \end{aligned}$$

Here $\mathbb{E}_{\mathcal{X}^d} \hat{G}_m(\mathcal{F}^d | \hat{h} \circ x^d)$ is Gaussian complexity of embeddings

$$\hat{h} \circ x^d = \{\hat{h}(x_1), \dots, \hat{h}(x_m) | x^d = x_1, \dots, x_m \in \mathcal{X}^d\}$$

where the expectation is over $\mathcal{X}^d$, and $\sigma^2 = \frac{1}{m} \sup_{f \in \mathcal{F}^d} \sum_{i=1}^m \text{Var}(\ell(f \circ \hat{h}(x_i^d), y_i^d))$ is the maximal variance over $\mathcal{F}^d$.

We remark that in Theorem 4.4, since we specialize to the case where $\mathcal{F}^p$ and $\mathcal{F}^d$ are sets of linear functions, we can replace the term $L^d \cdot \bar{G}_m(\mathcal{F}^d)$ in Theorem 4.2 by $(\sqrt{k'} \cdot \mathbb{E}_{\mathcal{X}^d} \hat{G}_m(\mathcal{F}^d | \hat{h} \circ x^d) + \sigma \sqrt{\log(1/\delta)/m})$ by utilizing the functional Bernstein inequality. This improvement can help us obtain Theorem 4.11. See Appendix A.2 for details.

Now we discuss the interpretation of Theorem 4.4. Typically, $\bar{G}_n(\mathcal{F}^p)$ is much smaller than $G_n(\mathcal{H})$ because $G_n(\mathcal{H})$ represents the complexity of the representation function, which is often complex. In practice, n is often large. Therefore, in the benign case where $\tilde{\nu} = \Theta(k)$ (when the condition number of α^p is $O(1)$), the dominating term will be $G_n(\mathcal{H})$. As we will show in the following subsections, this term typically scales as $O(\sqrt{1/n})$. Together, the theorem clearly shows when 1) the number of pre-training data is large, and 2) the least singular value of the last linear layer for pre-training is large, the transfer learning risk is small. On the other hand, if $\tilde{\nu}$ is small, then the bound becomes loose. This is consistent with prior counterexamples on multi-task pre-training (Maurer et al., 2016; Du et al., 2021) where the diversity is shown to be necessary.

4.3 What is diversity parameter for Linear Layers?

To prove Theorem 4.4, one of our key technical contributions is to show the following lemma that bridges the gap between Theorem 4.2 and Theorem 4.4

Lemma 4.5. *Under Assumption 3.1, 3.2, 4.3, we have*

$$d_{\mathcal{F}^d}(\hat{h}; h) \leq \frac{1}{\Omega(\tilde{\nu})} d_{\mathcal{F}^p, \mathcal{F}^p}(\hat{h}; h). \quad (4)$$

In intuition, it says that $\tilde{\nu}$ could serve as a lower bound for the diversity parameter ν . The take-away message is that we may wish to achieve a higher $\tilde{\nu}$ in order to increase the diversity of the pre-training models, thus improving its generality to various downstream tasks.

Technically, in proving the results we shall need to apply a *modified self-concordance* condition for better characterizing multinomial logistic regression (Bach et al., 2010). We note that the proofs of this part is very different from the multi-task setting studied in previous works (Du et al., 2021; Tripuraneni et al., 2020).

We define some additional notations for clarity and simplicity in this subsection. Let α' and α denote the parameters for $\hat{f}^p$ and f^p respectively. Let $\Phi(x) = \log(1 + \sum_{s=1}^{k-1} e^{x_s})$, for $x \in \mathbb{R}^{k-1}$, which is widely seen in multinomial regression tasks because the cross-entropy loss is inherently analogous to multinomial logistic loss.

In this subsection, we emphasize on the techniques required to show the following lemma, which incorporates the main difficulty in the proof for Lemma 4.5.

Lemma 4.6. *The Kullback-Leibler (KL) divergence between the true underlying conditional distribution of a multinomial logistic model and the distribution we obtained can be bounded from both sides with quadratic loss,*

$$\begin{aligned} & c_0 e^{-10q_0} \left\| \alpha'^\top \hat{h}(x^p) - \alpha^\top h(x^p) \right\|^2 \\ & \leq KL \left[\mathcal{P}(\cdot | \alpha^\top h(x^p)), \mathcal{P}(\cdot | \alpha'^\top \hat{h}(x^p)) \right] \\ & \leq \frac{1}{2} \left\| \alpha'^\top \hat{h}(x^p) - \alpha^\top h(x^p) \right\|^2, \end{aligned}$$

where $c_0 = \frac{1}{2} \lambda_{\min}(\Phi''(\alpha^\top h(x^p)))$ is the least eigenvalue of Hessian matrix for Φ , $q_0 = \max(\|\alpha'^\top \hat{h}(x^p)\|, \|\alpha^\top h(x^p)\|)$.

Remark 4.7. For the left hand side, the expression is related to the least eigenvalue of the Hessian matrix at $\alpha^\top h(x^p)$, whereas the least eigenvalue would depend on an unknown intermediate-term x' if we adopt Taylor's expansion.

Proof of Lemma 4.6. Below we use x for x^p for clarity. For generalized linear models,

$$\begin{aligned} KL \left[\mathcal{P}(\cdot | \alpha^\top h(x)), \mathcal{P}(\cdot | \alpha'^\top \hat{h}(x)) \right] &= \Phi(\alpha'^\top \hat{h}(x)) - \\ &\Phi(\alpha^\top h(x)) - \nabla \Phi(\alpha^\top h(x))^\top (\alpha'^\top \hat{h}(x) - \alpha^\top h(x)). \end{aligned}$$

Hence the divergence serves as the second-order remainder term according to Taylor's theorem.

For the right hand side, the gradient of $\Phi(x)$ at i^{th} -coordinate $\frac{\partial \Phi}{\partial x_i} = e^{x_i} / (1 + \sum_s e^{x_s})$, the Hessian matrix is

$$\frac{\partial^2 \Phi}{\partial x_i \partial x_j} = \begin{cases} \frac{e^{x_i} \cdot (1 + \sum_{s \neq i} e^{x_s})}{(1 + \sum_s e^{x_s})^2}, & i = j \\ \frac{-e^{x_i} e^{x_j}}{(1 + \sum_s e^{x_s})^2}, & i \neq j. \end{cases}$$

Let $\sigma(x) = \frac{1}{1 + \sum_s e^{x_s}} [e^{x_1}, \dots, e^{x_{k-1}}]^\top$, the Hessian matrix can be restated as

$$\nabla^2 \Phi = \text{diag}(\sigma(x)) - \sigma(x) \sigma(x)^\top.$$

For any non-zero vector y , we have

$$\begin{aligned} y^\top \nabla^2 \Phi y &= \sum_i \sigma(x)_i y_i^2 - (\sigma(x)^\top y)^2 \\ &\leq \max(\sigma(x)_i) \|y\|^2 \\ &\leq \|y\|^2 \end{aligned}$$

which implies its largest eigenvalue is no bigger than 1.

For the left hand side, it is very straightforward to see that the Hessian matrix is positive semi-definite. Though being non-negative, we point out that bounding the second-order

remainder terms from below with quadratic loss requires new techniques which we discuss below.

Since multinomial logistic regression is not strongly-convex, we need to find new techniques that would present benign properties to characterize the local landscape. Below we introduce a class of convex functions called *modified self-concordant* functions, which would be useful in quantitative analysis.

Definition 4.8 (Modified Self-concordance). Suppose $F: \mathbb{R}^p \mapsto \mathbb{R}$ is a three times differentiable convex function such that for some $R > 0$, for all $u, v \in \mathbb{R}^p$, the function $g: t \mapsto F(u + tv)$ satisfies for all $t \in \mathbb{R}$

$$|g'''(t)| \leq R \|v\|_2 \times g''(t) \quad (5)$$

Properties of self-concordance Self-concordance gives nice characterizations of local curvature of convex functions which plays important role in describing local convexity (Bach, 2014). Some useful results are given upon this condition (see (Bach et al., 2010, Proposition 1)), we list out the three main inequalities as below: For all $w, v \in \mathbb{R}^p, t \in \mathbb{R}$,

$$\begin{aligned} F(w + v) &\geq F(w) + vF'(w) + \\ &\quad \frac{v^\top F''(w)v}{R^2 \|v\|_2^2} \cdot \left(e^{-R\|v\|_2} + R\|v\|_2 - 1 \right), \\ F(w + v) &\leq F(w) + vF'(w) + \\ &\quad \frac{v^\top F''(w)v}{R^2 \|v\|_2^2} \cdot \left(e^{R\|v\|_2} - R\|v\|_2 - 1 \right), \\ e^{-tR\|v\|_2} F''(w) &\preceq F''(w + tv) \preceq e^{tR\|v\|_2} F''(w). \end{aligned}$$

The first two inequalities are refined characterizations of Taylor's expansion, while the last line presents bounds for Hessian matrix in the sense of positive semi-definiteness.

We find that multinomial logistic loss satisfies the *modified self-concordance* condition with $R = 5$.

Proposition 4.9. *For all $u, v \in \mathbb{R}^{k-1}$, the function $g: t \mapsto \Phi(u + tv)$ satisfies*

$$|g'''(t)| \leq 5 \|v\|_2 g''(t).$$

See Appendix A.2 for detailed derivations. Equipped with self-concordance, we are ready to give a lower bound of the divergence,

$$\begin{aligned} &\Phi(\alpha'^\top \hat{h}(x)) - \Phi(\alpha^\top h(x)) - \nabla \Phi(\alpha^\top h(x))^\top v \\ &\geq \frac{1}{2} v^\top e^{-5\|v\|_2} F''(\alpha^\top h(x)) v \\ &\geq \frac{1}{2} \lambda_{\min}(\Phi''(\alpha^\top h(x))) \|v\|^2 e^{-5\|v\|_2} \\ &\geq \frac{1}{2} \lambda_{\min}(\Phi''(\alpha^\top h(x))) \|v\|^2 e^{-5(\|\alpha'^\top \hat{h}(x)\| + \|\alpha^\top h(x)\|)} \\ &\geq \frac{1}{2} \lambda_{\min}(\Phi''(\alpha^\top h(x))) \|v\|^2 \exp(-10q_0) \end{aligned}$$

where $v = \alpha'^\top \hat{h}(x) - \alpha^\top h(x)$. This completes our proofs for Lemma 4.6. Please find the remaining details for completing the proof of Lemma 4.5 in Appendix A.2.

4.4 Linear Subspace Representation

Based on cross-entropy loss and linear predictors introduced in Section 4.2, we further assume the underlying representation is a projection onto a low-dimensional subspace. For $r \ll d$, let the representation be

$$\mathcal{H} = \{h|h(x) = B^\top x, B \in \mathbb{R}^{d \times r}\},$$

where B is a matrix with orthonormal columns. We require some additional regularity conditions. Following prior work (Du et al., 2021; Tripuraneni et al., 2020), we assume that $\|x\| \leq D$ and input data distribution satisfies the following condition.

Definition 4.10. The covariate distribution $P_x(\cdot)$ is Σ -subgaussian if for all $v \in \mathbb{R}^d$,

$$\mathbb{E}[\exp(v^\top x)] \leq \exp\left(\frac{\|\Sigma^{1/2}v\|^2}{2}\right)$$

where the covariance Σ further satisfies $\sigma_{\max}(\Sigma) \leq C$ and $\sigma_{\min}(\Sigma) \geq c > 0$ for universal constants c, C .

We have the following theorem that guarantees the performance of transfer learning.

Theorem 4.11. Suppose Assumption 3.1, 3.2, and 4.3 hold, data generation follows Definition 4.10. For a sufficiently large constant c_4 , we assume $n \geq c_4 d, m \geq c_4 r, D \leq c_4(\min(\sqrt{dr^2}, \sqrt{rm}))$. Then with probability at least $1 - \delta$, we have the Transfer Learning Risk upper bounded by:

$$O\left(\frac{1}{\tilde{\nu}} \left[\sqrt{k} \log(n) \left(\sqrt{\frac{kdr^2}{n}} + k\sqrt{\frac{r}{n}} \right) + \frac{k}{n^2} + \sqrt{\frac{\log(1/\delta)}{n}} \right] + (k')^{\frac{3}{2}} \sqrt{\frac{r}{m}} + k' \sqrt{\frac{\log(1/\delta)}{m}} \right)$$

To interpret this bound, consider the practically relevant scenario where $k' = O(1)$ (e.g., sentiment analysis), $m \ll n$, $k \ll n$ and $r \ll d$, in the benign case $\tilde{\nu} = \Omega(k)$, we have the transfer learning risk $\tilde{O}\left(\sqrt{dr^2/n} + \sqrt{r/m}\right)$. Note that this is exactly the desired theoretical guarantee because the first term accounts for using all pre-training data to learn the representation function and the second term accounts for using the downstream data to learn the last linear layer. This is significantly better than not using pre-training, in which case the risk scales $O\left(\sqrt{d/m}\right)$. Furthermore, for the linear representation learning setting, classic minimax bounds present a standard $\Omega(\sqrt{d/m})$ lower rate, which is also worse than our upper bound with representation learning (Foster et al., 2018; Abramovich and Grinshtein, 2018; Barnes and Ozgur, 2019).

4.5 Deep Neural Network Representation

In this subsection, we assume the underlying representation function to be a $\sigma = \tanh$ -activated neural network, which is often used in practice. Predictors are still required to be linear functions at the interest of NLP pre-training, i.e.,

$$\begin{aligned} \mathcal{H} &= \{h|h(x) = W_K \sigma(W_{K-1} \sigma(\cdots \sigma(W_1 x)))\}, \\ \mathcal{F}^p &= \{f|f(z) = \alpha^\top z, \alpha \in \mathbb{R}^{r \times (k-1)}, \\ &\quad \|\alpha_s\| \leq c_1 M(K)^2, s \in [k-1], \|\alpha^\top z\| \leq c_2\}, \\ \mathcal{F}^d &= \{f|f(z) = \alpha^\top z, \alpha \in \mathbb{R}^{r \times (k'-1)}, \\ &\quad \|\alpha_s\| \leq c_0 M(K)^2, s \in [k'-1], \|\alpha^\top z\| \leq c_3\}. \end{aligned}$$

Here M refer to constants that only depend on the network configuration, which satisfy: 1) for each $p \in [K]$, $\|W_p\|_{1,\infty} \leq M(p)$, and 2) $\|W_K\|_{\infty \rightarrow 2} \leq M(K)$.

Adapt Gaussian complexity results in Golowich et al. (2018) we have

$$\begin{aligned} G_n(\mathcal{H}) &\leq \tilde{O}\left(\frac{rM(K)^3 \cdot D\sqrt{K} \cdot \Pi_{p=1}^{K-1} M(p)}{\sqrt{n}}\right), \\ \bar{G}_n(\mathcal{F}^p|h \circ x^p) &\leq O\left(\frac{(k-1)M(K)^3}{\sqrt{n}}\right). \end{aligned}$$

Now we are ready to state our theorem for this practical setting of NLP pre-training.

Theorem 4.12. Under Assumption 3.1, 3.2, and 4.3, assume $M(K) \geq c_5$ for a universal constant c_5 . Then with probability at least $1 - \delta$, Transfer Learning Risk is upper bounded by

$$\begin{aligned} \tilde{O}\left(\frac{krM(K)^3 \cdot D\sqrt{K} \cdot \Pi_{p=1}^{K-1} M(p)}{\tilde{\nu}\sqrt{n}} + \frac{k^{\frac{3}{2}} M(K)^3}{\tilde{\nu}\sqrt{n}} + \frac{k'^{\frac{3}{2}} M(K)^3}{\sqrt{m}}\right). \end{aligned}$$

To interpret this bound, one can easily show that a standard supervised learning paradigm without pre-training would have a sample complexity of $\tilde{O}(krM(K)^3 \cdot D\sqrt{K} \cdot \Pi_{p=1}^{K-1} M(p)/\sqrt{m})$. Again, this theorem demonstrates: when 1) $n \gg m$ and 2) $\tilde{\nu}$ is large, the rate of transfer learning risk can be much smaller than that of the standard supervised learning algorithm.

5 CONCLUSION AND FUTURE WORK

This work theoretically prove the benefit of multi-class pre-training using the notion of class diversity. Our proof uses the vector-form Rademacher complexity chain rule and a modified self-concordance condition.

Future work First, our work is based on realizability assumptions (cf. Assumption 3.1 and 3.2) that are commonly adopted in transfer learning and classical PAC learning framework in order to present non-trivial statistical guarantees (Maurer et al., 2016; Du et al., 2021; Tripuraneni et al., 2020; Shalev-Shwartz and Ben-David, 2014). We believe our theorems can be extended to agnostic versions by relaxing these assumptions.

Second, if the target task is well-aligned with the source tasks, one can define more fine-grained notions to capture the task relevance. An example is (Chen et al., 2022), in which regression setting is studied. One interesting direction is extending their task relevance definition to the classification setting.

Finally, there has been some interesting recent work showing that one can do pre-training (i.e., masked word prediction) with the downstream dataset itself (which is usually smaller than typical pre-training corpora) and get good results (Krishna et al., 2022). Compared to the setting studied in this work, it might be harder to justify its performance through a “diversity” perspective because their settings are generally beyond the standard transfer learning scheme. Nevertheless, our interpretation of ν as a task-relatedness parameter might help shed light on these results, which is worthy of investigation.

Acknowledgements

This work was supported in part by NSF CCF 2212261, NSF IIS 2143493, NSF DMS-2134106, NSF CCF 2019844, NSF IIS 2110170, and a gift funding from Tencent.

References

- F. Abramovich and V. Grinshtein. High-dimensional classification by sparse logistic regression. *IEEE Transactions on Information Theory*, 65(5):3068–3079, 2018.
- F. Bach. Adaptivity of averaged stochastic gradient descent to local strong convexity for logistic regression. *The Journal of Machine Learning Research*, 15(1):595–627, 2014.
- F. Bach et al. Self-concordant analysis for logistic regression. *Electronic Journal of Statistics*, 4:384–414, 2010.
- Y. Bansal, G. Kaplun, and B. Barak. For self-supervised learning, Rationality implies generalization, provably. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=Srmggo3b3X6>.
- L. P. Barnes and A. Ozgur. Minimax bounds for distributed logistic regression. *arXiv preprint arXiv:1910.01625*, 2019.
- J. Baxter. A Model of Inductive Bias Learning. *Journal of artificial intelligence research*, 12:149–198, 2000.
- S. Boyd and L. Vandenberghe. *Convex Optimization*. Cambridge University Press, 2004. doi: 10.1017/CBO9780511804441.
- R. Caruana. Multitask Learning. *Machine learning*, 28(1): 41–75, 1997.
- Y. Chen, K. Jamieson, and S. Du. Active multi-task representation learning. In *International Conference on Machine Learning*, pages 3271–3298. PMLR, 2022.
- K. Clark, M.-T. Luong, Q. V. Le, and C. D. Manning. ELECTRA: Pre-training Text Encoders as Discriminators Rather Than Generators. In *International Conference on Learning Representations*, 2020. URL <https://openreview.net/forum?id=r1xMH1BtvB>.
- J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota, June 2019. Association for Computational Linguistics. doi: 10.18653/v1/N19-1423. URL <https://www.aclweb.org/anthology/N19-1423>.
- S. S. Du, W. Hu, S. M. Kakade, J. D. Lee, and Q. Lei. Few-Shot Learning via Learning the Representation, Provably. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=pW2Q2xLwIMD>.
- D. J. Foster, S. Kale, H. Luo, M. Mohri, and K. Sridharan. Logistic regression: The importance of being improper. In *Conference On Learning Theory*, pages 167–208. PMLR, 2018.
- N. Golowich, A. Rakhlin, and O. Shamir. Size-independent Sample Complexity of Neural Networks. In *Conference On Learning Theory*, pages 297–299. PMLR, 2018.
- K. Krishna, S. Garg, J. P. Bigham, and Z. C. Lipton. Downstream datasets make surprisingly good pretraining corpora. *arXiv preprint arXiv:2209.14389*, 2022.
- Z. Lan, M. Chen, S. Goodman, K. Gimpel, P. Sharma, and R. Soricut. ALBERT: A Lite BERT for Self-supervised Learning of Language Representations. In *International Conference on Learning Representations*, 2020. URL <https://openreview.net/forum?id=H1eA7AEtvS>.
- M. Ledoux and M. Talagrand. *Probability in Banach Spaces: Isoperimetry and Processes*. A Series of Modern Surveys in Mathematics Series. Springer, 1991. ISBN 9783540520139. URL <https://books.google.com.hk/books?id=cyKYDfvxRjsC>.
- J. D. Lee, Q. Lei, N. Saunshi, and J. Zhuo. Predicting what you already know helps: Provable self-supervised learn-

- ing. *Advances in Neural Information Processing Systems*, 34:309–323, 2021.
- Y. Liu, M. Ott, N. Goyal, J. Du, M. Joshi, D. Chen, O. Levy, M. Lewis, L. Zettlemoyer, and V. Stoyanov. RoBERTa: A Robustly Optimized BERT Pretraining Approach, 2020. URL <https://openreview.net/forum?id=SyxS0T4tvS>.
- P. Massart. About the Constants in Talagrand’s Concentration Inequalities for Empirical Processes. *Annals of Probability*, pages 863–884, 2000.
- A. Maurer. A vector-contraction inequality for Rademacher complexities. In *International Conference on Algorithmic Learning Theory*, pages 3–17. Springer, 2016.
- A. Maurer, M. Pontil, and B. Romera-Paredes. The Benefit of Multitask Representation Learning. *Journal of Machine Learning Research*, 17(81):1–32, 2016.
- A. Radford, K. Narasimhan, T. Salimans, and I. Sutskever. Improving Language Understanding by Generative Pre-Training. 2018.
- J. Robinson, S. Jegelka, and S. Sra. Strength from weakness: Fast learning using weak supervision. In *International Conference on Machine Learning*, pages 8127–8136. PMLR, 2020.
- N. Saunshi, O. Plevrakis, S. Arora, M. Khodak, and H. Khandeparkar. A Theoretical Analysis of Contrastive Unsupervised Representation Learning. In *International Conference on Machine Learning*, pages 5628–5637, 2019.
- S. Shalev-Shwartz and S. Ben-David. *Understanding machine learning: From theory to algorithms*. Cambridge university press, 2014.
- K. K. Thekumparampil, P. Jain, P. Netrapalli, and S. Oh. Sample Efficient Linear Meta-Learning by Alternating Minimization. *arXiv preprint arXiv:2105.08306*, 2021.
- C. Tosh, A. Krishnamurthy, and D. Hsu. Contrastive learning, multi-view redundancy, and linear models. In *Algorithmic Learning Theory*, pages 1179–1206. PMLR, 2021.
- N. Tripuraneni, M. Jordan, and C. Jin. On the theory of transfer learning: The importance of task diversity. *Advances in neural information processing systems*, 33: 7852–7862, 2020.
- N. Tripuraneni, C. Jin, and M. Jordan. Provable meta-learning of linear representations. In *International Conference on Machine Learning*, pages 10434–10443. PMLR, 2021.
- A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, Ł. Kaiser, and I. Polosukhin. Attention is all you need. *Advances in neural information processing systems*, 30, 2017.
- M. J. Wainwright. *High-dimensional statistics: A non-asymptotic viewpoint*, volume 48. Cambridge University Press, 2019.
- A. Wang, A. Singh, J. Michael, F. Hill, O. Levy, and S. R. Bowman. GLUE: A Multi-Task Benchmark and Analysis Platform for Natural Language Understanding. In *International Conference on Learning Representations*, 2019. URL <https://openreview.net/forum?id=rJ4km2R5t7>.
- Z. Yang, Z. Dai, Y. Yang, J. Carbonell, R. R. Salakhutdinov, and Q. V. Le. XLNet: Generalized autoregressive pre-training for language understanding. *Advances in neural information processing systems*, 32, 2019.
- Y. You, J. Li, S. Reddi, J. Hseu, S. Kumar, S. Bhojanapalli, X. Song, J. Demmel, K. Keutzer, and C.-J. Hsieh. Large batch optimization for deep learning: Training bert in 76 minutes. In *International Conference on Learning Representations*, 2020. URL <https://openreview.net/forum?id=Syx4wnEtvH>.
- Y. Zhang, A. Backurs, S. Bubeck, R. Eldan, S. Gunasekar, and T. Wagner. Unveiling Transformers with LEGO: a synthetic reasoning task. *arXiv preprint arXiv:2206.04301*, 2022.
- J. Y. Zou and R. P. Adams. Priors for Diversity in Generative Latent Variable Models. In *Proceedings of the 25th International Conference on Neural Information Processing Systems-Volume 2*, pages 2996–3004, 2012.

A TECHNICAL PROOFS

In Section 3, we have introduced Gaussian complexity. Let us restate for clarity.

Let μ be a probability distribution on a set $\mathcal{X} \subset \mathbb{R}^d$ and suppose that $x_1, \dots, x_n$ are independent samples selected according to μ . Let $\mathcal{Q}$ be a class of functions mapping from $\mathcal{X}$ to $\mathbb{R}^r$. Define random variable

$$\hat{G}_n(\mathcal{Q}) = \mathbb{E} \left[\sup_{q \in \mathcal{Q}} \frac{1}{n} \sum_{k=1}^r \sum_{i=1}^n g_{ki} q_k(x_i) \right] \quad (6)$$

as the empirical Rademacher complexity, where $q_k(x_i)$ is the k -th coordinate of the vector-valued function $q(x_i)$, g_{ki} ($k \in [r], i \in [n]$) are independent Gaussian $\mathcal{N}(0, 1)$ random variables. The Gaussian complexity of $\mathcal{Q}$ is $G_n(\mathcal{Q}) = E_\mu \hat{G}_n(\mathcal{Q})$.

Analogously to the above we can define the empirical Rademacher complexity for vector-valued functions as

$$\hat{R}_n(\mathcal{Q}) = \mathbb{E} \left[\sup_{q \in \mathcal{Q}} \frac{1}{N} \sum_{k=1}^r \sum_{i=1}^N \epsilon_{ki} q_k(x_i) \right] \quad (7)$$

where ϵ_{ki} ($k \in [r], i \in [n]$) are independent Rademacher $\text{Rad}(\frac{1}{2})$ random variables. Its population counterpart is defined as $R_n(\mathcal{Q}) = E_\mu[\hat{R}_n(\mathcal{Q})]$. Note that the superscripts existing in $\hat{G}$ and $\hat{R}$ imply that they are empirical measures.

A.1 Proofs for Section 4.1

We illustrate Theorem 4.2 in two stages. First we show pre-training representation difference can be upper bounded by constants and function class complexities. Then we transfer it to the downstream task through the diversity parameter.

Pre-training

Theorem A.1. *In pre-training, with probability at least $1 - \delta$, it holds that:*

$$\begin{aligned} d_{\mathcal{F}^p, f^p}(h'; h) &\leq 4\sqrt{\pi} L^p G_n(\mathcal{F}^p \circ \mathcal{H}) + 4B^p \sqrt{\frac{\log(2/\delta)}{n}} \\ &\leq 4096 L^p \left[\frac{\sqrt{k-1} D_{\mathcal{X}^p}}{n^2} + \log(n) [L(\mathcal{F}^p) G_n(\mathcal{H}) + \bar{G}_n(\mathcal{F}^p)] \right] + 4B^p \sqrt{\frac{\log(2/\delta)}{n}}. \end{aligned}$$

Proof. We begin with

$$d_{\mathcal{F}^p, f^p}(h'; h) \leq 2 \sup_{f \in \mathcal{F}^p, h \in \mathcal{H}} |R_p(f^p, h) - \hat{R}_p(f^p, h)|.$$

From the definition of Rademacher complexity (Wainwright, 2019, Theorem 4.12), with probability at least $1 - 2\delta$ we have

$$\sup_{f^p \in \mathcal{F}^p, h \in \mathcal{H}} |R_p(f^p, h) - \hat{R}_p(f^p, h)| \leq 2R_n(\ell(\mathcal{F}^p \circ \mathcal{H})) + 2B^p \sqrt{\frac{\log(1/\delta)}{n}}.$$

Next, we apply the vector contraction inequality (Maurer, 2016). For function class $\mathcal{F}$ whose output is in $\mathbb{R}^K$ with component $f_k(\cdot)$, and the function (h_i) s are some L -Lipschitz functions: $\mathbb{R}^K \mapsto \mathbb{R}$, we have

$$\mathbb{E}_\epsilon \sup_{f \in \mathcal{F}} \sum_{i=1}^n \epsilon_i h_i(f(x_i)) \leq \sqrt{2} L \mathbb{E}_\epsilon \sup_{f \in \mathcal{F}} \sum_{i=1}^n \sum_{k=1}^K \epsilon_{ik} f_k(x_i). \quad (8)$$

Hence for loss function ℓ satisfying $|\ell(x) - \ell(y)| \leq L\|x - y\|_2, \forall x, y \in \mathbb{R}^{k-1}$, the f takes value in $\mathbb{R}^{k-1}$ with component functions $f_s(\cdot), s \in [k-1]$, we have that population Rademacher complexity can be bounded by

$$R_n(\ell(\mathcal{F}^p \circ \mathcal{H})) = \mathbb{E}_{\mathcal{X}^p} \frac{1}{n} \mathbb{E}_\epsilon \sup_{f \in \mathcal{F}^p, h \in \mathcal{H}} \sum_{i=1}^n \epsilon_i \ell(f \circ h(x_i^p))$$

$$\begin{aligned}
 &\leq \mathbb{E}_{\mathcal{X}^p} \frac{1}{n} \sqrt{2} L^p \mathbb{E}_\epsilon \sup_{f \in \mathcal{F}^p, h \in \mathcal{H}} \sum_{i=1}^n \sum_{s=1}^{k-1} \epsilon_{is} f_s(h(x_i^p)) \\
 &= \sqrt{2} L^p R_n(\mathcal{F}^p \circ \mathcal{H}) \\
 &\leq \sqrt{\pi} L^p G_n(\mathcal{F}^p \circ \mathcal{H})
 \end{aligned}$$

where the last line uses the fact that Rademacher complexity is upper bounded by Gaussian complexity: $R_n(\mathcal{F}^p \circ \mathcal{H}) \leq \sqrt{\frac{\pi}{2}} G_n(\mathcal{F}^p \circ \mathcal{H})$. Therefore we have, with probability at least $1 - \delta$,

$$\begin{aligned}
 &d_{\mathcal{F}^p, \mathcal{F}^p}(h'; h) \\
 &\leq 4\sqrt{\pi} L^p G_n(\mathcal{F}^p \circ \mathcal{H}) + 4B^p \sqrt{\frac{\log(2/\delta)}{n}} \\
 &\leq 4096 L^p \left[\frac{\sqrt{k-1} D_{\mathcal{X}^p}}{n^2} + \log(n) [L(\mathcal{F}^p) G_n(\mathcal{H}) + \bar{G}_n(\mathcal{F}^p)] \right] + 4B^p \sqrt{\frac{\log(2/\delta)}{n}},
 \end{aligned}$$

where the last line uses decomposition of $G_n(\mathcal{F}^p \circ \mathcal{H})$ into the individual Gaussian complexities of $\mathcal{H}$ and $\mathcal{F}^p$, leverages an expectation version of novel chain rule for Gaussian complexities (Lemma A.2). $\square$

In the spirit of Gaussian complexity decomposition theorem (Tripuraneni et al., 2020, Theorem 7), we introduce the following decomposition result upon vector-form Gaussian complexities.

Lemma A.2. *We have the following vector form Gaussian complexity decomposition:*

$$\hat{G}_n(\mathcal{F}^p \circ \mathcal{H}) \leq \frac{8\sqrt{k-1} D_{\mathcal{X}^p}}{n^2} + 512C(\mathcal{F}^p \circ \mathcal{H}) \cdot \log(n) \quad (9)$$

where we use $C(\mathcal{F}^p \circ \mathcal{H}) = L(\mathcal{F}^p) \cdot \hat{G}_n(\mathcal{H}) + \bar{G}_n(\mathcal{F}^p)$ to represent the complexity measure of the composite function class.

Proof. Our proof extends (Tripuraneni et al., 2020, Theorem 7), which focuses on a multi-task scalar formulation. We further extend it to multi-class vector formulation. Specifically, on top of the representation class $\mathcal{H}$, they need to handle $\mathcal{F}^{\otimes t}$ (t is the number of tasks) while our objective is a single function class $\mathcal{F}^p$ of higher dimension ($\mathcal{F}^p$ is $(k-1)$ -dimensional for a k -class classification task). We note that our proof technique and that of previous works (Tripuraneni et al., 2020; Maurer et al., 2016) both hinge on several properties of Gaussian processes.

To bound the empirical composite function class $\mathcal{F}^p(\mathcal{H})$, note that vector-form Gaussian complexity is defined as

$$\begin{aligned}
 \hat{G}_n(\mathcal{F}^p \circ \mathcal{H}) &= \mathbb{E} \left[\frac{1}{n} \sup_{f(h) \in \mathcal{F}^p(\mathcal{H})} \sum_{s=1}^{k-1} \sum_{i=1}^n g_{is} f_s(h(x_i^p)) \right] \\
 &= \frac{1}{\sqrt{n}} \mathbb{E} \left[\sup_{f(h) \in \mathcal{F}^p(\mathcal{H})} Z_{f(h)} \right]
 \end{aligned}$$

where we define mean-zero process $Z_{f(h)} = \frac{1}{\sqrt{n}} \sum_{s=1}^{k-1} \sum_{i=1}^n g_{is} f_s(h(x_i^p))$, then $\mathbb{E} \sup_{f(h)} Z_{f(h)} = \mathbb{E} \sup_{f(h)} Z_{f(h)} - Z_{f'(h')} \leq \mathbb{E} \sup_{f(h), f'(h')} Z_{f(h)} - Z_{f'(h')}$. We further notice that $Z_{f(h)} - Z_{f'(h')}$ is a sub-gaussian random variable parameter

$$\begin{aligned}
 d^2(f(h), f'(h'))|x^p &= \frac{1}{n} \sum_{i=1}^n \|f(h(x_i^p)) - f'(h'(x_i^p))\|^2 \\
 &= \frac{1}{n} \sum_{s=1}^{k-1} \sum_{i=1}^n (f_s(h(x_i^p)) - f'_s(h'(x_i^p)))^2
 \end{aligned}$$

Dudley's entropy integral bound (Wainwright, 2019, Theorem 5.22) shows

$$\mathbb{E} \sup_{f(h), f'(h')} Z_{f(h)} - Z_{f'(h')}$$

$$\begin{aligned}
 &\leq 2\mathbb{E} \sup_{d(f(h), f'(h'))|x^{\mathcal{P}}| \leq \delta} Z_{f(h)} - Z_{f'(h')} + 32\mathcal{J}\left(\frac{\delta}{4}, D_{\mathcal{X}^{\mathcal{P}}}\right) \\
 &= 2\mathbb{E} \sup_{d(f(h), f'(h'))|x^{\mathcal{P}}| \leq \delta} Z_{f(h)} - Z_{f'(h')} + 32 \int_{\frac{\delta}{4}}^{D_{\mathcal{X}^{\mathcal{P}}}} \sqrt{\log N(u; \mathcal{F}^{\mathcal{P}}(\mathcal{H})|x^{\mathcal{P}})} du.
 \end{aligned}$$

It is straightforward to see the first term follows:

$$\mathbb{E} \sup_{d(f(h), f'(h'))|x^{\mathcal{P}}| \leq \delta} Z_{f(h)} - Z_{f'(h')} \leq \mathbb{E}[\|g\|]\delta \leq \sqrt{n(k-1)}\delta$$

We now turn to bound the second term by decomposing the distance metric into a distance over $\mathcal{F}^{\mathcal{P}}$ and a distance over $\mathcal{H}$. We claim that, for arbitrary $h \in \mathcal{H}$, $f \in \mathcal{F}^{\mathcal{P}}$, let h' be ϵ_1 -close to h in empirical l_2 -norm w.r.t inputs $x_1^{\mathcal{P}}, x_2^{\mathcal{P}} \dots, x_n^{\mathcal{P}}$. Given h' , let f' be ϵ_2 -close to f in empirical l_2 loss w.r.t $h'(x^{\mathcal{P}})$. Using the triangle inequality we have that

$$\begin{aligned}
 d(f(h), f'(h'))|x^{\mathcal{P}}| &= \sqrt{\frac{1}{n} \sum_{i=1}^n \|f(h(x_i^{\mathcal{P}})) - f'(h'(x_i^{\mathcal{P}}))\|^2} \\
 &\leq d(f(h), f(h'))|x^{\mathcal{P}}| + d(f(h'), f'(h'))|x^{\mathcal{P}}| \\
 &\leq \sqrt{\frac{1}{n} \sum_{i=1}^n \|f(h(x_i^{\mathcal{P}})) - f(h'(x_i^{\mathcal{P}}))\|^2} + \epsilon_2 \\
 &\leq L(\mathcal{F}^{\mathcal{P}}) \sqrt{\frac{1}{n} \sum_{i=1}^n \|h(x_i^{\mathcal{P}}) - h'(x_i^{\mathcal{P}})\|^2} + \epsilon_2 \\
 &= L(\mathcal{F}^{\mathcal{P}}) \cdot \epsilon_1 + \epsilon_2,
 \end{aligned}$$

where we have used that $\|f(x) - f(y)\| \leq L(\mathcal{F}^{\mathcal{P}})\|x - y\|$ for any $f \in \mathcal{F}^{\mathcal{P}}$.

As for the cardinality of the covering $C_{\mathcal{F}^{\mathcal{P}}(\mathcal{H})}$. Observe $|C_{\mathcal{F}^{\mathcal{P}}(\mathcal{H})}| = \sum_{h \in C_{\mathcal{H}(x^{\mathcal{P}})}} |C_{\mathcal{F}_h^{\mathcal{P}}}| \leq |C_{\mathcal{H}(x^{\mathcal{P}})}| \cdot \max_{h \in \mathcal{H}(x^{\mathcal{P}})} |C_{\mathcal{F}_h^{\mathcal{P}}}|$. This provides a bound on the metric entropy of

$$\log N(\epsilon_1 \cdot L(\mathcal{F}^{\mathcal{P}}) + \epsilon_2; \mathcal{F}^{\mathcal{P}}(\mathcal{H})|x^{\mathcal{P}}) \leq \log N(\epsilon_1; \mathcal{H}|x^{\mathcal{P}}) + \max_{h(x^{\mathcal{P}})} \log N(\epsilon_2; \mathcal{F}^{\mathcal{P}}|h \circ x^{\mathcal{P}}).$$

Applying the covering number upper bound with $\epsilon_1 = \frac{\epsilon}{2 \cdot L(\mathcal{F}^{\mathcal{P}})}$, $\epsilon_2 = \frac{\epsilon}{2}$ gives a bound of entropy integral of a ,

$$\begin{aligned}
 &\int_{\frac{\delta}{4}}^{D_{\mathcal{X}^{\mathcal{P}}}} \sqrt{\log N(u; \mathcal{F}^{\mathcal{P}}(\mathcal{H})|x^{\mathcal{P}})} du \\
 &\leq \int_{\frac{\delta}{4}}^{D_{\mathcal{X}^{\mathcal{P}}}} \sqrt{\log N\left(\frac{u}{2L(\mathcal{F}^{\mathcal{P}})}; \mathcal{H}|x^{\mathcal{P}}\right)} du + \int_{\frac{\delta}{4}}^{D_{\mathcal{X}^{\mathcal{P}}}} \max_{h \circ x^{\mathcal{P}}} \sqrt{\log N\left(\frac{u}{2}; \mathcal{F}^{\mathcal{P}}|h \circ x^{\mathcal{P}}\right)} du
 \end{aligned}$$

From the Sudakov minoration theorem (Wainwright, 2019, Theorem 5.30) for Gaussian processes and the fact that packing numbers at scale u upper bounds the covering number at scale $\forall u > 0$ we find:

$$\log N(u; \mathcal{H}|x^{\mathcal{P}}) \leq 4 \left(\frac{\sqrt{n} \hat{G}_n(\mathcal{H})}{u} \right)^2, \quad \log N(u; \mathcal{F}^{\mathcal{P}}|h(x^{\mathcal{P}})) \leq 4 \left(\frac{\sqrt{n} \hat{G}_n(\mathcal{F}^{\mathcal{P}}|h \circ x^{\mathcal{P}})}{u} \right)^2.$$

Combining the definition of worst-case Gaussian complexity with the aforementioned results we have

$$\hat{G}_n(\mathcal{F}^{\mathcal{P}} \circ \mathcal{H}) \leq 2\sqrt{k-1}\delta + 256 \log \frac{4D_{\mathcal{X}^{\mathcal{P}}}}{\delta} \left(L(\mathcal{F}^{\mathcal{P}}) \hat{G}_n(\mathcal{H}) + \bar{G}_n(\mathcal{F}^{\mathcal{P}}) \right),$$

substitute δ with $\frac{4D_{\mathcal{X}^{\mathcal{P}}}}{n^2}$, proof is completed. $\square$

Downstream learning Next we turn to the second stage and come up with theoretical guarantees by using inexact $\hat{h}$ learned from the first stage.

Theorem A.3. *In the downstream task, we have that with probability at least $1 - \delta$,*

$$R_d(\hat{f}^d, \hat{h}) - R_d(f^d, h) \leq d_{\mathcal{F}^d}(\hat{h}; h) + 4\sqrt{\pi}L^d \cdot \bar{G}_m(\mathcal{F}^d) + 4B^d \sqrt{\frac{\log(2/\delta)}{m}}$$

Proof. Assumption 3.1 implies

$$\mathbb{E}_{x^d, y^d} [\ell(g^d(x^d), y^d)] = R_d(f^d, h).$$

To start, let $\tilde{f} = \arg \min_{f \in \mathcal{F}^d} R_d(f, \hat{h})$ and $R_d(\hat{f}^d, \hat{h}) - R_d(f^d, h)$ equals

$$\left[R_d(\tilde{f}, \hat{h}) - R_d(f^d, h) \right] + \left[R_d(\hat{f}^d, \hat{h}) - R_d(\tilde{f}, \hat{h}) \right]$$

where the first term satisfies

$$\begin{aligned} & \inf_{\tilde{f} \in \mathcal{F}^d} \left[R_d(\tilde{f}, \hat{h}) - R_d(f^d, h) \right] \\ & \leq \sup_{f^d \in \mathcal{F}^d} \inf_{\tilde{f} \in \mathcal{F}^d} \left[R_d(\tilde{f}, \hat{h}) - R_d(f^d, h) \right] \\ & = d_{\mathcal{F}^d}(\hat{h}, h) \end{aligned}$$

The second term follows the similar lines of Theorem A.1

$$R_d(\hat{f}^d, \hat{h}) - R_d(\tilde{f}, \hat{h}) \leq 4\sqrt{\pi}L^d \mathbb{E}_{\mathcal{X}^d} \hat{G}_m(\mathcal{F}^d | \hat{h} \circ x^d) + 4B^d \sqrt{\frac{\log(1/\delta)}{m}}$$

Again we make use of the worst-case argument

$$\mathbb{E}_{\mathcal{X}^d} \hat{G}_m(\mathcal{F}^d | \hat{h} \circ x^d) \leq \bar{G}_m(\mathcal{F}^d).$$

Combining the results gives the statement. $\square$

Proof of main Theorem 4.2 Having introduced class diversity parameter, proof is directly completed via combination of Theorem A.1 and Theorem A.3.

A.2 Proofs for Section 4.2

We could provide a better dependence on the boundedness noise parameters in Theorem A.3 using Bernstein inequality. We present the following corollary which has data-dependence in the Gaussian complexities.

Corollary A.4. *Presuming Assumption 3.1 holds, we have that then with probability at least $1 - \delta$,*

$$\begin{aligned} & R_d(\hat{f}^d, \hat{h}) - R_d(f^d, h) \\ & \leq d_{\mathcal{F}^d}(\hat{h}; h) + 4\sqrt{\pi}L^d \cdot \mathbb{E}_{\mathcal{X}^d} \hat{G}_m(\mathcal{F}^d | \hat{h} \circ x^d) + 4\sigma \sqrt{\frac{\log(2/\delta)}{m}} + 50B^d \frac{\log(2/\delta)}{m} \end{aligned}$$

Proof. Denote $Z = \sup_f |\hat{R}_d(f, \hat{h}) - R_d(f, h)|$, we apply the functional Bernstein inequality (Massart, 2000, Theorem 3) to control the fluctuations. With probability at least $1 - \delta$, we have

$$Z \leq 2\mathbb{E}[Z] + 4\frac{\sigma}{\sqrt{m}} \sqrt{\log\left(\frac{1}{\delta}\right)} + 35\frac{B^d}{m} \log\left(\frac{1}{\delta}\right), \quad (10)$$

where $\sigma^2 = \frac{1}{m} \sup_f \sum_{i=1}^m \text{Var}(\ell(f \circ \hat{h}(x_i^d), y_i^d))$. Thus

$$\begin{aligned} \mathbb{E}[Z] & \leq 2\mathbb{E}_{\mathcal{X}^d} \hat{R}_m(l(\mathcal{F}^d) | \hat{h} \circ x^d) \\ & \leq 2\mathbb{E}_{\mathcal{X}^d} \sqrt{2}L^d \hat{R}_m(\mathcal{F}^d | \hat{h} \circ x^d) \\ & \leq 2\mathbb{E}_{\mathcal{X}^d} \sqrt{\pi}L^d \hat{G}_m(\mathcal{F}^d | \hat{h} \circ x^d), \end{aligned}$$

where the second line uses vector-based contraction principle, the last line upper bounds the empirical Rademacher complexity by Gaussian counterparts. $\square$

Proof of Theorem 4.4 Observe that

$$\ell(\eta; y) = -y^\top \eta + \log \left(1 + \sum_{s=1}^{k-1} e^{\eta_s} \right), \ell(\eta; y) \leq \|\eta\|$$

and

$$\left| \frac{\partial \ell(\eta; y)}{\partial \eta_i} \right| = \left| y_i - \frac{e^{\eta_i}}{1 + \sum_{s=1}^{k-1} e^{\eta_s}} \right|,$$

$$|\nabla_\eta \ell(\eta; y)| \leq \sqrt{k-1},$$

so it is $L^P = \sqrt{k-1}$ -Lipschitz. By definition the class $\mathcal{F}^P$ with parameters $\|\alpha_s\|_2 \leq O(1), s \in [k-1]$, we obtain that $L(\mathcal{F}^P) = O(\sqrt{k-1})$ since for any $x, y \in \mathbb{R}^r$, any $f \in \mathcal{F}^P$ we have

$$\begin{aligned} \|f(x) - f(y)\|^2 &= \|\alpha^\top x - \alpha^\top y\|^2 \\ &\leq \sum_{s=1}^{k-1} (\langle \alpha_s, x - y \rangle)^2 \\ &\leq \sum_{s=1}^{k-1} \|\alpha_s\|^2 \|x - y\|^2 \\ &\leq c_1^2 (k-1) \|x - y\|^2 \end{aligned}$$

In conclusion we have

- Pre-training loss $\ell(\cdot, y^P)$ is $\sqrt{k-1}$ -Lipschitz.
- Downstream loss $\ell(\cdot, y^d)$ is $\sqrt{k'-1}$ -Lipschitz.
- Linear layer f is $L(\mathcal{F}^P) = O(\sqrt{k-1})$ -Lipschitz.

Consider task-specific function classes for characterizing class-diversity parameters. From Lemma 4.6 and Lemma 4.5 we know that

$$\nu = \Omega(\tilde{\nu}), \quad \tilde{\nu} = \sigma_r(\alpha_1 \alpha_1^\top).$$

Combining these pieces of results then the proof is completed.

With the following proposition, we interpret the cross-entropy loss in the well-specified model under our multinomial logistic model distribution.

Proposition A.5. *Under Assumption 3.2, for the cross entropy loss ℓ we have*

$$\begin{aligned} \mathbb{E}_{y \sim \mathcal{P}(\cdot | f \circ h(x))} [\ell(\hat{f} \circ \hat{h}(x), y)] - \ell(f \circ h(x), y) &= KL \left[\mathcal{P}(\cdot | f \circ h(x)), \mathcal{P}(\cdot | \hat{f} \circ \hat{h}(x)) \right] \\ &= KL \left[\mathcal{P}(\cdot | \alpha^\top h(x)), \mathcal{P}(\cdot | \alpha'^\top \hat{h}(x)) \right]. \end{aligned}$$

Recall that α' and α are parameters for $\hat{f}$ and f respectively. The proof is straightforward by applying Assumption 3.2.

Proof of Proposition 4.9

Proof. Let $P(t; v^0) = 1 + \sum_s e^{u_s + t v_s}$ and $P(t; v^i) = \sum_s v_s^i e^{u_s + t v_s}, i > 1$. Then we use multinomials P to represent derivatives of $g(t)$

$$g(t) = \log(P(T; V^0))$$

$$g'(t) = \frac{P(t; v^1)}{P(t; v^0)}$$

$$g''(t) = \frac{P(t; v^2)P(t; v^0)}{P(t; v^0)^2}$$

$$g'''(t) = \frac{P(t; v^3)P(t; v^0)^2 - 3P(t; v^2)P(t; v^1)P(t; v^0) + 2P(t; v^1)^3}{P(t; v^0)^3}$$

Let $r_s = e^{u_s + v_s t}$, hence

$$g''(t) = \frac{(\sum_s v_s^2 r_s) \cdot (1 + \sum_s r_s) - (\sum_s v_s r_s)^2}{(1 + \sum_s r_s)^2}$$

$$= \frac{\sum_{i < j} r_i r_j (v_i - v_j)^2 + \sum_i v_i^2 r_i}{(1 + \sum_s r_s)^2}$$

In the following we expand $g'''(t)$ as:

$$\frac{\sum_{i < j} r_i r_j (v_i - v_j)^2 [\sum_k (v_i + v_j - 2v_k) r_k] + \sum_i v_i^3 r_i + \sum_i \sum_j v_i^2 r_i r_j (2v_i - 3v_j)}{(1 + \sum_s r_s)^3}$$

$$= \frac{\sum_{i < j} r_i r_j (v_i - v_j)^2 [\sum_k (v_i + v_j - 2v_k) r_k] + \sum_i v_i^2 r_i (v_i (1 + 2 \sum_j r_j) - 3 \sum_j v_j r_j)}{(1 + \sum_s r_s)^3},$$

observe that

$$\frac{1}{1 + \sum_s r_s} \left| \sum_k (v_i + v_j - 2v_k) r_k \right| \leq \sum_k |v_i + v_j - 2v_k| \frac{r_k}{1 + \sum_s r_s} \leq 4\|v\|_2$$

$$\frac{1}{1 + \sum_s r_s} \left| v_i (1 + 2 \sum_j r_j) - 3 \sum_j v_j r_j \right| \leq 5\|v\|_2$$

Substitute these into definition of *self-concordance* then proof is completed. $\square$

Now we are ready to give a lower bound of KL -divergence,

$$\begin{aligned} & \Phi(\alpha'^\top \hat{h}(x)) - \Phi(\alpha^\top h(x)) - \nabla \Phi(\alpha^\top h(x))^\top v \\ & \geq \frac{1}{2} v^\top e^{-5\|v\|_2} F''(\alpha^\top h(x)) v \\ & \geq \frac{1}{2} \lambda_{\min}(\Phi''(\alpha^\top h(x))) \|v\|^2 e^{-5\|v\|_2} \\ & \geq \frac{1}{2} \lambda_{\min}(\Phi''(\alpha^\top h(x))) \|v\|^2 e^{-5(\|\alpha'^\top \hat{h}(x)\| + \|\alpha^\top h(x)\|)} \\ & \geq \frac{1}{2} \lambda_{\min}(\Phi''(\alpha^\top h(x))) \|v\|^2 \exp(-10q_0) \end{aligned}$$

where $v = \alpha'^\top \hat{h}(x) - \alpha^\top h(x)$. Proof for Lemma 4.6 is completed. $\square$

Proof of Lemma 4.5

Proof. For function classes $\mathcal{F}^p$, $\mathcal{F}^d$ and data samples generated from multinomial logistic regression distribution (see Assumption 3.2), the worst-case representation difference is similar to that in multi-task analysis (Tripuraneni et al., 2020, Lemma 1):

$$d_{\mathcal{F}^d}(\hat{h}; h) = \sup_{f^d \in \mathcal{F}^d} \inf_{f' \in \mathcal{F}^d} \mathbb{E} \left\{ \ell(f' \circ \hat{h}(x^d), y^d) - \ell(f^d \circ h(x^d), y^d) \right\}$$

$$\leq \sup_{\|\alpha_s\| \leq c_0} \inf_{\|\alpha'_s\| \leq c_0} \frac{1}{2} \mathbb{E}_{\mathcal{X}^d} \left\| \alpha'^\top \hat{h}(x^d) - \alpha^\top h(x^d) \right\|^2, \quad \text{here } s \in [k' - 1]$$

$$\begin{aligned}
 &= \sum_{s=1}^{k'-1} \sup_{\|\alpha_s\| \leq c_0} \inf_{\|\alpha'_s\| \leq c_0} \frac{1}{2} \mathbb{E}_{\mathcal{X}^d} \left(\alpha'_s{}^\top \hat{h}(x^d) - \alpha_s{}^\top h(x^d) \right)^2 \\
 &\leq (k' - 1) \frac{c_0^2}{2} \sigma_1(\Lambda_{sc}(\hat{h}, h)).
 \end{aligned}$$

The first line is because of Proposition A.5 and Lemma 4.6. In the last line, the inner infima is considered as the partial minimization of a convex quadratic form (see (Boyd and Vandenberghe, 2004, Example 3.15, Appendix A.5.4)).

Define population covariance if representations $\hat{h}$ and h as

$$\Lambda(\hat{h}, h) = \begin{bmatrix} \mathbb{E}[\hat{h}(x)\hat{h}(x)^\top] & \mathbb{E}[\hat{h}(x)h(x)^\top] \\ \mathbb{E}[h(x)\hat{h}(x)^\top] & \mathbb{E}[h(x)h(x)^\top] \end{bmatrix} = \begin{bmatrix} F_{\hat{h}\hat{h}} & F_{\hat{h}h} \\ F_{h\hat{h}} & F_{hh} \end{bmatrix}$$

$\Lambda_{sc}(\hat{h}, h) = F_{hh} - F_{h\hat{h}}(F_{\hat{h}\hat{h}})^\dagger F_{\hat{h}h}$ is the generalized Schur complement of h with respect to $\hat{h}$.

Next we control the pre-training representation difference, which is subtler,

$$\begin{aligned}
 &d_{\mathcal{F}^p, f^p}(\hat{h}; h) \\
 &\geq \inf_{\alpha'} c_0 \mathbb{E}_{\mathcal{X}^p} \left[\exp(-10 \max(\|\alpha'^\top \hat{h}(x^p)\|, \|\alpha^\top h(x^p)\|)) \cdot \left\| \alpha'^\top \hat{h}(x^p) - \alpha^\top h(x^p) \right\|^2 \right],
 \end{aligned}$$

which is because of Proposition A.5 and Lemma 4.6.

It is known

$$\begin{aligned}
 &\mathbb{E}_{\mathcal{X}^p} \left[\exp(-10 \max(\|\alpha'^\top \hat{h}(x^p)\|, \|\alpha^\top h(x^p)\|)) \cdot \left\| \alpha'^\top \hat{h}(x^p) - \alpha^\top h(x^p) \right\|^2 \right] \\
 &\geq e^{-10c_2} \left\| \alpha'^\top \hat{h}(x^p) - \alpha^\top h(x^p) \right\|^2
 \end{aligned}$$

Hence this metric could be claimed to be lower bounded as,

$$\begin{aligned}
 &\Omega \left(\inf_{\alpha'} \mathbb{E}_{x^p} \left\| \alpha'^\top \hat{h}(x^p) - \alpha^\top h(x^p) \right\|^2 \right) \\
 &= \Omega \left(\alpha_1^\top \Lambda_{sc}(\hat{h}; h) \alpha_1 \right) \\
 &= \Omega \left(\text{tr}(\Lambda_{sc}(\hat{h}; h) C) \right), \quad \text{where } C = \alpha_1 \alpha_1^\top.
 \end{aligned}$$

In the second line, we redefine α_1 as parameter α of pre-training for clarity. In this way we conclude that,

$$d_{\mathcal{F}^p, f^p}(\hat{h}; h) = \Omega \left(\text{tr}(\Lambda_{sc}(\hat{h}, h) C) \right) = \Omega \left(\sigma_1(\Lambda_{sc}(\hat{h}, h)) \sigma_r(C) \right),$$

where C implies expansion of representation $h(x) \in \mathbb{R}^r$, and its condition number $\sigma_r(C)$ indicates how spread out this vector is in $\mathbb{R}^r$:

$$C = \sum_{s=1}^{k-1} (\alpha_1)_s (\alpha_1)_s^\top = \alpha_1 \alpha_1^\top, \quad \alpha_1 \in \mathbb{R}^{r \times (k-1)}$$

Aforementioned calculations show

$$d_{\mathcal{F}^d}(\hat{h}; h) \leq \frac{1}{\Omega(\tilde{\nu})} d_{\mathcal{F}^p, f^p}(\hat{h}; h), \quad \tilde{\nu} = \sigma_r(C).$$

Proof is completed. $\square$

A.3 Proofs for Section 4.4

Proof. We begin with bounding each of the complexity terms in the Corrolary A.4.

We make use of data-dependent inequalities (Tripuraneni et al., 2020, Lemma 4) to help upper bound related quantities. Intuitively Definition 4.10 implies tail-bound properties in a sub-gaussian process.

•

$$\begin{aligned}\hat{G}_n(\mathcal{H}) &= \frac{1}{n} \mathbb{E} \left[\sup_{B \in \mathcal{H}} \sum_{k=1}^r \sum_{i=1}^n g_{ki} b_k^\top x_i^p \right] \\ &= O \left(\sqrt{\frac{dr^2}{n}} \right)\end{aligned}$$

•

$$\begin{aligned}\hat{G}_n(\mathcal{F}^p | h \circ x^p) &= \frac{1}{n} \mathbb{E} \left[\sup_{\alpha_1, \dots, \alpha_{k-1}} \sum_{s=1}^{k-1} \sum_{i=1}^n g_{is} \alpha_s^\top B^\top x_i^p \right] \\ &= \frac{c_1(k-1)}{n} \mathbb{E} \left\| \sum_{i=1}^n g_{is} B^\top x_i^p \right\| \\ &= \frac{c_1(k-1)}{\sqrt{n}} \sqrt{\text{tr}(B^\top \Sigma B)}\end{aligned}$$

then $\bar{G}_n(\mathcal{F}^p) \leq O((k-1)\sqrt{\frac{r}{n}})$.

• Similarly,

$$\hat{G}_m(\mathcal{F}^d | h \circ x^d) \leq \frac{c_1(k'-1)}{\sqrt{m}} \sqrt{\sum_{i=1}^r \sigma_i(\hat{B}^\top \Sigma \hat{B})}$$

then $\bar{G}_m(\mathcal{F}^d) \leq O((k'-1)\sqrt{\frac{r}{m}})$.

- boundedness parameter $D_{\mathcal{X}^p} = \sup_{\alpha, B} \|\alpha^\top B^\top x\| = c_2$
- cross entropy $\ell(\eta; y) = -y^\top \eta + \log(1 + \sum_{s=1}^{k-1} e^{\eta_s})$, then $\left| \frac{\partial \ell(\eta; y)}{\partial \eta_i} \right| = \left| y_i - \frac{e^{\eta_i}}{1 + \sum_{s=1}^{k-1} e^{\eta_s}} \right|$, $|\nabla_\eta \ell(\eta; y)| \leq \sqrt{k-1}$, so it is $L^p = \sqrt{k-1}$ -Lipschitz in its first coordinate uniformly over its second for pre-training and $L^d = \sqrt{k'-1}$ -Lipschitz for downstream task.
- $|\ell(\eta; y)| \leq O(\|\eta\|)$, where $\|\eta\| = \|x^\top B^p \alpha\| \leq c_2$.

In Corollary A.4, we define and compute the maximal variance term σ^2 as,

$$\begin{aligned}\sigma^2 &= \frac{1}{m} \sup_{f^d \in \mathcal{F}^d} \sum_{i=1}^m \text{Var}(\ell'(f^d \circ \hat{h}(x_i^d), y_i^d)) \\ &\leq \frac{k'-1}{m} \sup_{f^d \in \mathcal{F}^d} \sum_{i=1}^m \text{Var}(f^d \circ \hat{h}(x_i^d)) \\ &= \frac{k'-1}{m} \sup_{\|\alpha_s\| \leq O(1)} \sum_{s=1}^{k'-1} \sum_{i=1}^m \text{Var}(\alpha_s^\top \hat{B}^\top x_i^d) \\ &= \frac{(k'-1)^2}{m} \sup_{\|\alpha_s\| \leq O(1)} \sum_{i=1}^m (\alpha_s \hat{B})^\top \Sigma \hat{B} \alpha_s \\ &= (k'-1)^2 O(\|\hat{B} \Sigma \hat{B}\|_2) \\ &= O((k'-1)^2)\end{aligned}$$

With these results in hand, we are now prepared to apply Corollary A.4, w.p. at least $1 - \delta$

$$R_d(\hat{f}^d, \hat{h}) - R_d(f^d, h)$$

$$\leq d_{\mathcal{F}^d}(\hat{h}; h) + 4\sqrt{\pi}L^d \cdot \mathbb{E}_{\mathcal{X}^d} \hat{G}_m \left(\mathcal{F}^d \middle| \hat{h} \circ x^d \right) + 4\sigma \sqrt{\frac{\log(2/\delta)}{m}} + 50B^d \frac{\log(2/\delta)}{m}$$

where $\hat{G}_m(\mathcal{F}^d | \hat{h} \circ x^d)$ is defined in Theorem 4.4.

Thus $L^d \cdot \mathbb{E}_{\mathcal{X}^d} \hat{G}_m \left(\mathcal{F}^d \middle| \hat{h} \circ x^d \right) \leq L^d \bar{G}_m(\mathcal{F}^d) \leq O((k' - 1)^{\frac{3}{2}} \sqrt{\frac{r}{m}})$, $\sigma \leq O(k' - 1)$, and $B^d \leq O(\sqrt{k' - 1}D)$. Further, we obtain upper bound of worst-case representation difference by diversity parameter and adoption of Theorem A.1: w.p. at least $1 - \delta$

$$\begin{aligned} & d_{\mathcal{F}^d}(\hat{h}; h) \\ & \leq \frac{d_{\mathcal{F}^d, \mathcal{F}^p}(\hat{h}; h)}{\nu} \\ & \leq \frac{1}{\nu} \left\{ 4096L \left[\log(n) \cdot [L(\mathcal{F}^p) \cdot G_n(\mathcal{H}) + \bar{G}_n(\mathcal{F}^p)] + \frac{\sqrt{k-1}D_{\mathcal{X}^p}}{n^2} \right] + 4B \sqrt{\frac{\log(2/\delta)}{n}} \right\} \\ & \lesssim \frac{1}{\nu} \left\{ \sqrt{k} \log(n) \left(\sqrt{\frac{kdr^2}{n}} + k\sqrt{\frac{r}{n}} \right) + \frac{k}{n^2} + \sqrt{\frac{\log(1/\delta)}{n}} \right\} \end{aligned}$$

The last thing to consider for completing the proof for Theorem 4.11 is giving accurate characterization of diversity parameter ν , which we leave for the next subsection. $\square$

A.4 Proofs for Section 4.5

Proof. In deep neural network, we first review complexity quantities. Adapted from Theorem 8 (Golowich et al., 2018), we have

$$\begin{aligned} \hat{R}_n(\mathcal{N}) & \leq \left(\frac{2}{n} \Pi_{p=1}^K M(p) \right) \sqrt{(K+1+\log d) \cdot \max_{j \in [d]} \sum_{i=1}^n x_{i,j}^2} \\ & \leq \frac{2D\sqrt{K+1+\log d} \cdot \Pi_{p=1}^K M(p)}{\sqrt{n}}. \end{aligned}$$

where $x_{i,j}$ denotes the j -th coordinate of vector x_i .

Then we proceed to bound the Gaussian complexities for our deep neural network and prove Theorem 4.12. Recall that under the conditions of the result we can use former results to verify the task diversity condition is satisfied with parameters $\Omega(\tilde{\nu})$ with $\tilde{\nu} = \sigma_r(\alpha_1 \alpha_1^\top) > 0$. We can see that $\|\mathbb{E}_x[\hat{h}(x)h^*(x)^\top]\|_2 \leq \mathbb{E}_x\|\hat{h}(x)h^*(x)\| \leq O(M(K)^2)$ using the norm bound from. Hence under this setting we can choose c_1 sufficiently large so that $c_1 M(K)^2 \gtrsim \frac{M(K)^2}{c} c_2$. The condition $M(K) \gtrsim 1$ in the theorem statement is simply used to clean up the final bound.

In order to instantiate Theorem 4.2 we begin by bounding each of the complexity terms in the expression.

- For the feature learning complexity in the training phase, we leverage above results, then

$$\begin{aligned} \hat{G}_n(\mathcal{H}) & = \frac{1}{n} \mathbb{E}_{\mathcal{W}_K} \left[\sup_{k=1}^r \sum_{i=1}^n g_{ki} h_k(x_i^p) \right] \leq \sum_{k=1}^r \hat{G}_n(h_k(x_i^p)) \\ & \leq \log(n) \cdot \sum_{k=1}^r \hat{R}_n h_k(x_i^p) \leq r \log(n) \frac{2D\sqrt{K+1+\log d} \cdot \Pi_{p=1}^K M(p)}{\sqrt{n}}. \end{aligned}$$

This also implies the population Gaussian complexity.

- By definition the class $\mathcal{F}$ as linear maps with parameters $\|\alpha_s\|_2 \leq c_1 M(K)^2, \forall s \in [k-1]$, we obtain that $L(\mathcal{F}) = c_1 \sqrt{k-1} M(K)^2$.
- For the complexity of learning $\mathcal{F}^p$ in the training phase we obtain,

$$\hat{G}_n(\mathcal{F}^p | h \circ x^p) = \frac{1}{n} \mathbb{E}_g \left[\sup_{\alpha \in \mathcal{F}} \sum_{s=1}^{k-1} \sum_{i=1}^n g_{is} \alpha_s^\top h(x_i^p) \right] \lesssim \frac{(k-1)M(K)^2}{n} \mathbb{E}_g \left[\left\| \sum_{i=1}^n g_{is} h(x_{1i}) \right\| \right]$$

Table 1: **Performance of diversity-regularized BERT pre-training with different values of diversity factor λ .** We finetune the pre-trained model on 8 downstream tasks from GLUE benchmark and evaluate them on their dev sets. All results are “mean (std)” from 5 runs with different random seeds. For MNLI, we average the accuracies on its matched and mismatched dev sets. For MRPC and QQP, we average their accuracy and F1 scores. For STS-B, we average Pearson’s correlation and Spearman’s correlation. All other tasks uses accuracy as the metric. The better-than-baseline numbers are underlined, and the best numbers are highlighted in boldface.

Model	MNLI	MRPC	SST-2	CoLA	QQP	QNLI	RTE	STS-B
BERT-base ($\lambda = 0.005$)	<u>84.17</u> (0.23)	<u>87.16</u> (1.81)	92.48 (0.19)	59.99 (0.28)	<u>89.42</u> (0.08)	<u>88.11</u> (0.54)	<u>67.28</u> (3.43)	<u>89.33</u> (0.07)
BERT-base ($\lambda = 0.05$)	<u>84.01</u> (0.10)	<u>86.35</u> (5.15)	<u>93.00</u> (0.16)	<u>62.66</u> (1.07)	<u>89.46</u> (0.03)	87.64 (0.44)	60.64 (6.08)	<u>89.57</u> (0.13)
BERT-base ($\lambda = 0.5$)	<u>84.00</u> (0.20)	<u>89.42</u> (0.51)	<u>92.93</u> (0.24)	60.76 (0.71)	<u>89.33</u> (0.12)	88.01 (0.23)	<u>67.93</u> (1.18)	<u>89.22</u> (0.23)
BERT-base (reproduced)	83.96 (0.08)	86.14 (4.64)	92.64 (0.20)	61.46 (0.74)	89.28 (0.09)	88.10 (0.27)	63.64 (6.64)	89.19 (0.07)

$$\lesssim \frac{(k-1)M(K)^2}{n} \sqrt{\sum_{i=1}^n \|h(x_i^p)\|^2} \lesssim \frac{(k-1)M(K)^2}{\sqrt{n}} \max_i \|h(x_i^p)\|.$$

For $\tanh$ activation function, we simply have

$$\|h(x)\|^2 = \|W_K r_{K-1}\|_2^2 \leq \|W_K\|_{\infty \rightarrow 2}^2,$$

where r_{K-1} denotes ourput of the $K-1$ th layer,

$$\|h(x)\| \leq O(M(K)).$$

In conclusion we obtain

$$\bar{G}_n(\mathcal{F}^p) \leq O\left(\frac{(k-1)M(K)^3}{\sqrt{n}}\right).$$

- Similarly

$$\hat{G}_m(\mathcal{F}^d | h \circ x^d) \leq O\left(\frac{(k'-1)M(K)^3}{\sqrt{m}}\right)$$

Then for **Regularity conditions** we have

- Boundedness parameter $D_{\mathcal{X}^p} = \sup_{\alpha, h} \|\alpha^\top h(x^p)\| = c_2$.
- Pre-training loss is $L^p = \sqrt{k-1}$ -Lipschitz and $B^p = c_2$ -bounded.
- Downstream loss is $L^d = \sqrt{k'-1}$ -Lipschitz and $B^d = c_3$ -bounded.

Assembling the previous complexity arguments shows the transfer learning risk is bounded by

$$\begin{aligned} &\lesssim \frac{L^p}{\tilde{\nu}} \left(\log(n) \left[L(\mathcal{F}^p) r \log(n) \frac{D\sqrt{K}\Pi_{p=1}^K M(p)}{\sqrt{n}} + \frac{kM(K)^3}{\sqrt{n}} \right] \right) + \frac{L^d k' M(K)^3}{\sqrt{m}} \\ &+ \left(\frac{1}{\tilde{\nu}} \max \left(\frac{L^p \sqrt{k} D_{\mathcal{X}^p}}{n^2}, B^p \sqrt{\frac{\log(1/\delta)}{n}} \right) + B^d \sqrt{\frac{\log(1/\delta)}{m}} \right) \end{aligned}$$

Substitute regularity conditions into it, then the risk is simplified as stated in Theorem 4.12. $\square$

B EXPERIMENTS

Our theoretical analysis in previous sections implies that the diversity of the model parameter matrix at the linear output layer in pre-training has a significant impact on the transfer capability, in the sense that the larger ν (diversity parameter of f^p), the smaller the risk. Therefore, we could *explicitly add a diversity regularizer* to the linear output layer to increase

diversity. Motivated by this, we propose to add the following diversity regularizer to the original BERT pre-training loss so that it becomes:

$$L'(\Theta) = L(\Theta) - \lambda \cdot \ln \det(\alpha^P (\alpha^P)^\top), \quad (11)$$

where Θ denotes the set of all model parameters, λ is a hyper-parameter that controls the magnitude of the diversity regularization, $\det(\cdot)$ denote the determinant of a matrix, and α^P is the model parameter matrix at the output linear layer. This type of diversity regularizer was proposed in Zou and Adams (2012). This regularization technique is different from prior work because it is specifically designed for multi-class pre-training: we only add the diversity regularizer to the last linear layer.

We use the above diversity-regularized loss (along with the original ℓ_2 -regularization) to pretrain BERT-base models under different values of diversity factor λ . Then we fine-tune them on 7 classification tasks and 1 regression task from the GLUE benchmark (Wang et al., 2019) to evaluate their transfer performance.⁵ Our pre-training and finetuning implementations are based on the opensource code released by Nvidia.⁶ We use the same pre-training data as the original BERT (i.e., English Wikipedia + TorontoBookCorpus).⁷ Our detailed pre-training and finetuning hyper-parameters along with other experimental details are reported in Appendix B.1.

In Table B, we report our performance on the dev sets of the 8 downstream tasks. All the experiments are repeated 5 times with different random seeds, and we report their mean values along with the standard deviations. The complete experiment results (including full MNLI, QQP, and MRPC results) can be found in Appendix B.1. From Table B, we note that adding the diversity regularization could generally improve the performance on these downstream tasks. In particular, when $\lambda = 0.5$, our pre-trained model outperforms the original BERT-base on 6 out of 8 tasks (with 3 of them being significant), while achieving comparable performance on the other 2 tasks. Although our model is slightly behind the original BERT on CoLA and QNLI, such a performance gap is not statistically significant. Besides, we also see that our model with $\lambda = 0.5$ achieves a much more stable performance (i.e., smaller std) on tasks with scarce finetuning data ($< 4K$ samples in MRPC and RTE). Our results, albeit still preliminary, demonstrate the potential of such a simple diversity-regularizer. It could be an effective and simple performance booster for any of the existing pre-trained NLP models (e.g., XLNet (Yang et al., 2019), RoBERTa (Liu et al., 2020), ALBERT (Lan et al., 2020), etc) with negligible computation and implementation cost. We leave the development of the more advanced diversity regularizer as a future work.

B.1 More Details

Full statistics (including matched and mismatched dev sets for MNLI, accuracy and F1 scores for MRPC and QQP, and (Pearson’s correlation + Spearman’s correlation)/2 for STS-B. All other tasks uses accuracy as the metric) could be found in Table B.1.

Table 2: Full statistics on GLUE dev sets.

Model	Statistics	MNLI(m/mm)	MRPC(acc/F1)	SST-2	CoLA	QQP(acc/F1)	QNLI	RTE	STS-B
$\lambda = 0.005$	mean	83.96/84.37	84.90/89.42	92.48	59.99	90.96/87.88	88.11	67.28	89.33
	std	0.26/0.21	2.28/1.34	0.19	0.28	0.05/0.11	0.54	3.43	0.07
$\lambda = 0.05$	mean	83.88/84.14	83.72/88.98	93.00	62.66	90.97/87.96	87.64	60.64	89.57
	std	0.04/0.16	6.69/3.62	0.16	1.07	0.05/0.04	0.44	6.08	0.13
$\lambda = 0.5$	mean	83.96/84.04	87.75/91.09	92.93	60.76	90.85/87.81	88.01	67.93	89.22
	std	0.15/0.24	0.52/0.50	0.24	0.71	0.10/0.14	0.23	1.18	0.23
BERT-base	mean	83.85/84.07	83.48/88.80	92.64	61.46	90.87/87.68	88.10	63.64	89.19
	std	0.13/0.04	6.08/3.19	0.20	0.74	0.07/0.11	0.27	6.64	0.07

⁵We do not report the WNLI (classification) task due to its reported issues of the task in Devlin et al. (2019).

⁶Distributed under Apache License: <https://github.com/NVIDIA/DeepLearningExamples/tree/master/PyTorch/LanguageModeling/BERT>

⁷Collected and pre-processed using the code and script included in the open-source code: <https://github.com/NVIDIA/DeepLearningExamples/tree/master/PyTorch/LanguageModeling/BERT>

Complete statistics Here we provide complete results on GLUE dev sets over 5 random seeds.

Table 3: Performance of reproduced BERT-base model.

	GLUE	MNLI(m/mm)	MRPC(acc/F1)	SST-2	CoLA	QQP(acc/F1)	QNLI	RTE	STS-B
seed	42	83.90/84.04	71.32/82.44	92.66	60.11	90.94/87.72	87.58	66.79	89.25
	0	83.86/84.12	86.27/90.18	92.43	62.05	90.74/87.53	88.19	54.29	89.07
	393	83.78/84.06	86.76/90.63	92.54	61.42	90.93/87.84	88.29	70.36	89.19
	78	84.05/84.02	86.76/90.63	92.55	61.50	90.89/87.60	88.12	57.14	89.18
	3837	83.66/84.11	86.27/90.13	93.00	62.20	90.87/87.73	88.33	69.64	89.26
	mean	83.85/84.07	83.48/88.80	92.64	61.46	90.87/87.68	88.10	63.64	89.19
	std	0.13/0.04	6.08/3.19	0.20	0.74	0.07/0.11	0.27	6.64	0.07

Table 4: Performance of $\lambda = 0.005$ regularized pre-training model.

	GLUE	MNLI(m/mm)	MRPC(acc/F1)	SST-2	CoLA	QQP(acc/F1)	QNLI	RTE	STS-B
seed	42	84.24/84.43	87.25/90.72	92.20	59.99	90.94/87.85	87.09	67.14	89.40
	0	83.91/83.96	86.27/90.47	92.43	60.06	91.03/87.88	88.24	70.71	89.36
	393	83.84/84.51	85.54/89.52	92.55	59.48	90.89/87.68	88.33	71.07	89.22
	78	84.23/84.44	84.80/89.45	92.78	60.13	90.95/87.97	88.71	65.71	89.38
	3837	83.56/84.53	80.64/86.93	92.43	60.30	91.01/88.00	88.17	61.79	89.28
	mean	83.96/84.37	84.90/89.42	92.48	59.99	90.96/87.88	88.11	67.28	89.33
	std	0.26/0.21	2.28/1.34	0.19	0.28	0.05/0.11	0.54	3.43	0.07

Table 5: Performance of $\lambda = 0.05$ regularized pre-training model.

	GLUE	MNLI(m/mm)	MRPC(acc/F1)	SST-2	CoLA	QQP(acc/F1)	QNLI	RTE	STS-B
seed	42	83.88/84.22	86.52/90.27	92.89	61.12	91.06/87.95	86.93	65.00	89.52
	0	83.83/83.92	88.97/92.00	93.00	64.36	90.96/87.93	87.73	54.29	89.65
	393	83.96/83.98	70.59/81.92	93.12	62.22	90.94/88.03	87.47	61.79	89.65
	78	83.86/84.26	87.50/91.06	92.78	63.13	90.94/87.97	88.26	68.93	89.69
	3837	83.86/84.30	85.04/89.66	93.23	62.49	90.93/87.91	87.82	53.21	89.33
	mean	83.88/84.14	83.72/88.98	93.00	62.66	90.97/87.96	87.64	60.64	89.57
	std	0.04/0.16	6.69/3.62	0.16	1.07	0.05/0.04	0.44	6.08	0.13

Table 6: Performance of $\lambda = 0.5$ regularized pre-training model.

	GLUE	MNLI(m/mm)	MRPC(acc/F1)	SST-2	CoLA	QQP(acc/F1)	QNLI	RTE	STS-B
seed	42	83.75/83.87	87.75/90.89	93.00	59.79	90.88/87.75	87.58	65.71	89.00
	0	83.84/84.30	88.24/91.56	92.66	60.94	90.87/87.77	88.14	67.86	89.25
	393	83.98/83.68	87.99/91.46	93.12	60.99	90.98/88.04	88.22	68.21	89.19
	78	84.02/84.29	87.99/91.33	93.23	60.22	90.87/87.86	88.12	68.93	89.65
	3837	84.19/84.05	86.76/90.21	92.66	61.86	90.67/87.61	87.98	68.93	89.03
	mean	83.96/84.04	87.75/91.09	92.93	60.76	90.85/87.81	88.01	67.93	89.22
	std	0.15/0.24	0.52/0.50	0.24	0.71	0.10/0.14	0.23	1.18	0.23

Finally, we report detailed hyperparameter settings below.

Pre-training Hyperparameters for pre-training are shown in Table 7.

Hyperparam	phase-1	phase-2
Number of Layers	12	12
Hidden size	768	768
FFN inner hidden size	3072	3072
Attention heads	12	12
Steps	7038	1563
Optimizer	LAMB	LAMB
Learning Rate	9e-3	6e-3
β_1	0.9	0.9
β_2	0.999	0.999
WarmUp	28.43 %	12.80 %
Batch Size	65536	32768

Table 7: Hyperparameters used in pre-training our models. We use the LAMB optimizer (You et al., 2020) for large-batch pretraining of the BERT model, where β_1 and β_2 are its two hyper-parameters.

Finetuning Hyperparameters for downstream tasks are shown in Table 8. We adapt these hyperparameters from Liu et al. (2020), Devlin et al. (2019), and Yang et al. (2019).

	LR	BSZ	# EP	WARMUP	WD	FP16	SEQ
CoLA	1.00E-05	32	20	6%	0.1	O2	128
SST-2	3.00E-05	32	10	6%	0.1	O2	128
MNLI	3.00E-05	32	5	6%	0.1	O2	128
QNLI	3.00E-05	32	10	6%	0.1	O2	128
QQP	3.00E-05	32	5	6%	0.1	O2	128
RTE	3.00E-05	16	5	6%	0.1	O2	128
MRPC	3.00E-05	16	5	6%	0.1	O2	128

Table 8: The hyperparameters used in finetuning our model in downstream tasks. LR: learning rate. BSZ: batch size. #EP: number of epochs. WARMUP: warmup ratio. FP16: automatic mixed precision (AMP) level. SEQ: input sequence length.

Computing infrastructure We pretrain our (diversity-regularized) BERT-base models using 32 Nvidia V100 GPUs (32GB RAM each), and the finetuning of the model uses 4 Nvidia V100 GPUs.