
The Power and Limitation of Pretraining-Finetuning for Linear Regression under Covariate Shift

Jingfeng Wu*

Johns Hopkins University
uuujf@jhu.edu

Difan Zou*

The University of Hong Kong
dzou@cs.hku.hk

Vladimir Braverman

Johns Hopkins University
vova@cs.jhu.edu

Quanquan Gu

University of California, Los Angeles
qgu@cs.ucla.edu

Sham M. Kakade

Harvard University
sham@seas.harvard.edu

Abstract

We study linear regression under covariate shift, where the marginal distribution over the input covariates differs in the source and the target domains, while the conditional distribution of the output given the input covariates is similar across the two domains. We investigate a transfer learning approach with pretraining on the source data and finetuning based on the target data (both conducted by online SGD) for this problem. We establish sharp instance-dependent excess risk upper and lower bounds for this approach. Our bounds suggest that for a large class of linear regression instances, transfer learning with $O(N^2)$ source data (and scarce or no target data) is as effective as supervised learning with N target data. In addition, we show that finetuning, even with only a small amount of target data, could drastically reduce the amount of source data required by pretraining. Our theory sheds light on the effectiveness and limitation of pretraining as well as the benefits of finetuning for tackling covariate shift problems.

1 Introduction

In *transfer learning* [Pan and Yang, 2009, Sugiyama and Kawanabe, 2012], an algorithm is provided with abundant data from a source domain and scarce or no data from a target domain, and aims to train a model that generalizes well on the target domain. A simple yet effective approach is to *pretrain* a model with the rich source data and then *finetune* the model with the available target data via, e.g., *stochastic gradient descent* (SGD) (see, e.g., Yosinski et al. [2014]). Despite its wide applicability in practice, the power and limitation of the pretraining-finetuning based transfer learning framework is not fully understood in theory. The focus of this work is to consider this issue in a specific transfer learning setup known as *covariate shift* [Pan and Yang, 2009, Sugiyama and Kawanabe, 2012], where the source and target distributions differ in their marginal distributions over the input, but coincide in their conditional distribution of the output given the input.

Regarding the theory of learning with covariate shift, there exists a rich set of results [Ben-David et al., 2010, Germain et al., 2013, Mansour et al., 2009, Mohri and Muñoz Medina, 2012, Cortes and Mohri, 2014, Cortes et al., 2019, Kpotufe and Martinet, 2018, Hanneke and Kpotufe, 2019, Ma et al., 2022] for the (regularized) empirical risk minimizer, which minimizes the empirical loss over the source data and target data (if available) with potential regularization terms (e.g., ℓ_2 -regularization). However, in most of these works [Ben-David et al., 2010, Germain et al., 2013, Mansour et al., 2009, Mohri and Muñoz Medina, 2012, Cortes and Mohri, 2014, Cortes et al., 2019], the generalization

*Equal Contribution

error on the target domain is bounded by the sum of a vanishing term (e.g., the training error) and a divergence between the two domains (see, e.g., discussions in Kpotufe and Martinet [2018], with a few notable exceptions that we will discuss later). Such bounds are very pessimistic because the additive error contributed by the source-target divergence only captures the *worst case* performance gap caused by distribution mismatch [David et al., 2010] and is too crude to describe the intriguing properties of pretraining-finetuning across different domains.

In this paper, we take a different approach to directly study the generalization performance of the pretraining-finetuning method. In particular, we consider linear regression under covariate shift, and an online SGD estimator which is firstly trained with the source data and then finetuned with the target data. We derive a target domain risk bound that is stated as a function of (i) the spectrum of the source and target population data covariance matrices, (ii) the amount of source and target data, and (iii) the (initial) stepsizes for pretraining and finetuning (see Theorem 3.1 for more details). Moreover, a nearly matching lower bound is provided to justify the tightness of our upper bound. The derived bounds comprehensively characterize the effects of pretraining and finetuning for *each* covariate shift problem and *each* algorithm configuration, based on which we make the following important observations:

- We compare the generalization performance (i.e., target domain excess risk) of pretraining (with source data) vs. supervised learning (with target data). We show that, for a large class of problems, $O(N^2)$ source data is sufficient for pretraining to match the performance of supervised learning with N target data.
- We next show the benefits of finetuning with scarce target data. In particular, for the problem class considered before, finetuning can reduce by at least constant factors the amount of source data required by pretraining. Moreover, there exist problem instances for which the pretraining-finetuning approach requires *polynomially* less amount of total data than pretraining (with source data) or supervised learning (with target data).
- Finally, our bounds can also be applied to the supervised learning setting, i.e., linear regression with last iterate SGD. In this case, our upper bound sharpens that of Wu et al. [2021] by a logarithmic factor, and as a consequence we close the gap between the upper and lower bounds for last iterate SGD when the signal-to-noise ratio is bounded.

Notation. For two positive-value functions $f(x)$ and $g(x)$ we write $f(x) \lesssim g(x)$ or $f(x) \gtrsim g(x)$ if $f(x) \leq cg(x)$ or $f(x) \geq cg(x)$ for some absolute constant $c > 0$ respectively, and we write $f(x) \approx g(x)$ if $f(x) \lesssim g(x) \lesssim f(x)$. For two vectors $\mathbf{u}$ and $\mathbf{v}$ in a Hilbert space, their inner product is denoted by $\langle \mathbf{u}, \mathbf{v} \rangle$ or equivalently, $\mathbf{u}^\top \mathbf{v}$. For a matrix $\mathbf{A}$, its spectral norm is denoted by $\|\mathbf{A}\|_2$. For two matrices $\mathbf{A}$ and $\mathbf{B}$ of appropriate dimension, their inner product is defined as $\langle \mathbf{A}, \mathbf{B} \rangle := \text{tr}(\mathbf{A}^\top \mathbf{B})$. For a positive semi-definite (PSD) matrix $\mathbf{A}$ and a vector $\mathbf{v}$ of appropriate dimension, we write $\|\mathbf{v}\|_{\mathbf{A}}^2 := \mathbf{v}^\top \mathbf{A} \mathbf{v}$. For a symmetric matrix $\mathbf{A}$ and a PSD matrix $\mathbf{B}$, we write $\|\mathbf{A}\|_{\mathbf{B}}^2 := \|\mathbf{B}^{-\frac{1}{2}} \mathbf{A} \mathbf{B}^{-\frac{1}{2}}\|_2^2$. The Kronecker/tensor product is denoted by $\otimes$. For a set $\mathbb{S}$, we use $|\mathbb{S}|$ to denote its cardinality.

1.1 Additional Related Work

We review some additional works that are mostly related to ours.

Learning under Covariate Shift. Kpotufe and Martinet [2018], Pathak et al. [2022] proposed new similarity measures to the source and target domains, and proved covariate shift bounds that do not contain an additive error of the divergence between the source and target distribution. Compared to our results, theirs can be applied to nonlinear regression/classifications as well; however in the case of linear regression, our bounds are more fine-grained and are tight upto constant factors for a broad class of problems (see Theorem 3.2), beyond being only optimal in the worst case.

It is worth noting that Hanneke and Kpotufe [2019] studied the value of target data in addressing covariate shift problems. Their discussion is based on the minimax risk bounds afforded by a given number of source and target data. In contrast, our discussion on the benefits of finetuning with target data is based on a completely different perspective, which is by comparing the *sample inflation* [Bahadur, 1967, 1971, Zou et al., 2021a] between pretraining-finetuning vs. pretraining vs. supervised learning, i.e., for *each* covariate shift problem instance, how much source (and target) data

are necessary for pretraining (and finetuning) to match the performance of supervised learning with certain amount of target data.

More recently, Ma et al. [2022] studied covariate shift problem in the nonparametric kernel regression setting, with the assumption that the density ratio (or second moment ratio) between the target and source distribution is bounded. Their results are similar to ours in that their bounds reflect the effect of the spectrum of the source population data covariance. Since our results are dimension-free, our bounds can also be applied in the nonparametric kernel regression setting. There are two notable differences: firstly, their estimator is (weighted) ridge regression and ours is given by SGD; moreover, our results do not rely on the bounded density ratio or bounded second moment condition.

In addition, there is a vast literature on constructing more sample-efficient transfer learning algorithms, e.g., importance weighting methods [Shimodaira, 2000, Cortes et al., 2010] and learning invariant representations [Arjovsky et al., 2019, Wu et al., 2019], to mention a few. Along this line, Lei et al. [2021] proposed nearly minimax optimal estimator for linear regression under distribution shift, but their method relies on the knowledge of target population covariance matrix. Developing new transfer learning algorithms is beyond the agenda in this paper.

SGD. The pretraining and finetuning discussed in this work are both conducted by online SGD, therefore our results are closely related to the generalization analysis of online SGD for linear regression in the supervised learning context [Bach and Moulines, 2013, Dieuleveut et al., 2017, Jain et al., 2017a,b, Ge et al., 2019, Zou et al., 2021b, Varre et al., 2021, Wu et al., 2021]. From a technical point of view, our theoretical results can be viewed as an extension of the SGD analysis from the supervised learning setting to the covariate shift setting.

2 Problem Setup

Transfer Learning. We use $\mathbf{x}$ to denote a covariate in a Hilbert space (that can be d -dimensional or countably infinite dimensional), and $y \in \mathbb{R}$ to denote its response. Consider a source and a target data distribution, denoted by $\mathcal{D}_{\text{source}}$ and $\mathcal{D}_{\text{target}}$ respectively. In the problem of *transfer learning*, we are given with M data sampled independently from the source distribution, and N data sampled independently from the target distribution (where $N \ll M$ or even $N = 0$), denoted by

$$(\mathbf{x}_i, y_i)_{i=1}^{M+N}, \text{ where } (\mathbf{x}_i, y_i) \sim \begin{cases} \mathcal{D}_{\text{source}}, & i = 1, \dots, M; \\ \mathcal{D}_{\text{target}}, & i = M + 1, \dots, M + N. \end{cases}$$

The goal of transfer learning is to learn a model based on the $M + N$ data that can generalize on the target domain. We are particularly interested in the *covariate shift* problem in transfer learning, where the source and target distributions satisfy: $\mathcal{D}_{\text{source}}(y|\mathbf{x}) = \mathcal{D}_{\text{target}}(y|\mathbf{x})$ but $\mathcal{D}_{\text{source}}(\mathbf{x}) \neq \mathcal{D}_{\text{target}}(\mathbf{x})$.

Linear Regression under Covariate Shift. A covariate shift problem is formally defined in the context of linear regression by Definitions 1 and 2.

Definition 1 (Covariances conditions). Assume that each entry and the trace of the source and target data covariance matrices are finite. Denote the source and target data covariance matrices by

$$\mathbf{G} := \mathbb{E}_{\mathcal{D}_{\text{source}}}[\mathbf{x}\mathbf{x}^\top], \quad \mathbf{H} := \mathbb{E}_{\mathcal{D}_{\text{target}}}[\mathbf{x}\mathbf{x}^\top],$$

respectively, and denote their eigenvalues by $(\mu_i)_{i \geq 1}$ and $(\lambda_i)_{i \geq 1}$, respectively. For convenience assume that both $\mathbf{G}$ and $\mathbf{H}$ are strictly positive definite.

Definition 2 (Model conditions). For a parameter $\mathbf{w}$, define its source and target risks by

$$\text{Risk}_{\text{source}}(\mathbf{w}) := \frac{1}{2} \mathbb{E}_{\mathcal{D}_{\text{source}}}(y - \mathbf{w}^\top \mathbf{x})^2, \quad \text{Risk}_{\text{target}}(\mathbf{w}) := \frac{1}{2} \mathbb{E}_{\mathcal{D}_{\text{target}}}(y - \mathbf{w}^\top \mathbf{x})^2,$$

respectively. Assume that there is a parameter $\mathbf{w}^*$ that *simultaneously* minimizes both source and target risks, i.e., $\mathbf{w}^* \in \arg \min_{\mathbf{w}} \text{Risk}_{\text{source}}(\mathbf{w}) \cap \arg \min_{\mathbf{w}} \text{Risk}_{\text{target}}(\mathbf{w})$. For convenience assume that $\mathbf{w}^*$ is unique.

We remark that the strict positive definiteness of $\mathbf{G}$ and $\mathbf{H}$ in Definition 1 and the uniqueness of $\mathbf{w}^*$ in Definition 2 are only made for the ease of presentation. Otherwise one can set $\mathbf{w}^*$ to be the minimum-norm solution, i.e., $\mathbf{w}^* = \arg \min\{\|\mathbf{w}\|_2 : \mathbf{w} \in \arg \min_{\mathbf{w}} \text{Risk}_{\text{source}}(\mathbf{w}) \cap \arg \min_{\mathbf{w}} \text{Risk}_{\text{target}}(\mathbf{w})\}$.

$\arg \min_{\mathbf{w}} \text{Risk}_{\text{target}}(\mathbf{w})$, and our results still hold. This argument also holds in a reproducing kernel Hilbert space [Schölkopf et al., 2002].

Our definitions of linear regression under covariate shift follow from Lei et al. [2021], Ma et al. [2022]. In the literature, covariate shift problem often assumes the source and target distributions differ in their marginal distributions over the input, but coincide in their conditional distribution of the output given the input (see, e.g., Sections 1.3.3 and 1.3.4 in Sugiyama et al. [2012]). When applied to the well-specified linear regression models, i.e., $\mathbb{E}[y|\mathbf{x}] = \mathbf{x}^\top \mathbf{w}^*$ for some parameter $\mathbf{w}^*$, the latter condition turns out to require the optimal parameter $\mathbf{w}^*$ is identical for both source and target domains. Therefore, the problems described by Definitions 1 and 2 include at least all well-specified linear regression problems under standard covariate shift conditions.

Excess Risk. For linear regression under covariate shift, the performance of a parameter $\mathbf{w}$ is measured by its *target domain excess risk*, i.e.,

$$\text{ExcessRisk}(\mathbf{w}) := \text{Risk}_{\text{target}}(\mathbf{w}) - \text{Risk}_{\text{target}}(\mathbf{w}^*) = \frac{1}{2} \langle \mathbf{H}, (\mathbf{w} - \mathbf{w}^*) \otimes (\mathbf{w} - \mathbf{w}^*) \rangle.$$

SGD. The transfer learning algorithm of our interests is pretraining-finetuning via *online stochastic gradient descent with geometrically decaying stepsizes*² (SGD). Without loss of generality, we assume the SGD iterates are initialized from $\mathbf{w}_0 = \mathbf{0}$. Then the SGD iterates are sequentially updated as follows:

$$\begin{aligned} \mathbf{w}_t &= \mathbf{w}_{t-1} - \gamma_{t-1}(\mathbf{x}_t \mathbf{x}_t^\top \mathbf{w}_{t-1} - \mathbf{x}_t y_t), \quad t = 1, \dots, M + N, \\ \text{where } \gamma_t &= \begin{cases} \gamma_0/2^\ell, & 0 \leq t < M, \ell = \lfloor t/\log(M) \rfloor; \\ \gamma_M/2^\ell, & M \leq t < N, \ell = \lfloor (t-M)/\log(N) \rfloor, \end{cases} \end{aligned} \quad (\text{SGD})$$

and the output is the last iterate, i.e., $\mathbf{w}_{M+N}$. Here γ_0 and γ_M are two hyperparameters that correspond to the initial stepsizes for pretraining and finetuning, respectively. In both pretraining and finetuning phases, the stepsize scheduler in (SGD) is epoch-wisely a constant and decays geometrically every certain number of epochs, which is widely used in deep learning [He et al., 2015]. We note that such (SGD) for linear regression has been analyzed by Ge et al. [2019], Wu et al. [2021] in the context of supervised learning. Our goal in this work is to understand the generalization of (SGD) in the covariate shift problems.

Assumptions. The following assumptions [Zou et al., 2021b, Wu et al., 2021] are crucial in our analysis.

Assumption 1 (Fourth moment conditions). *Assume that for both source and target distribution the fourth moment of the covariates is finite. Moreover:*

A *There is a constant $\alpha > 0$ such that for every PSD matrix $\mathbf{A}$ it holds that*

$$\mathbb{E}_{\mathcal{D}_{\text{source}}}[\mathbf{x}\mathbf{x}^\top \mathbf{A}\mathbf{x}\mathbf{x}^\top] \preceq \alpha \cdot \text{tr}(\mathbf{G}\mathbf{A}) \cdot \mathbf{G}, \quad \mathbb{E}_{\mathcal{D}_{\text{target}}}[\mathbf{x}\mathbf{x}^\top \mathbf{A}\mathbf{x}\mathbf{x}^\top] \preceq \alpha \cdot \text{tr}(\mathbf{H}\mathbf{A}) \cdot \mathbf{H}.$$

Clearly, it must hold that $\alpha \geq 1$.

B *There is a constant $\beta > 0$ such that for every PSD matrix $\mathbf{A}$ it holds that*

$$\mathbb{E}_{\mathcal{D}_{\text{source}}}[\mathbf{x}\mathbf{x}^\top \mathbf{A}\mathbf{x}\mathbf{x}^\top] - \mathbf{G}\mathbf{A}\mathbf{G} \succeq \beta \cdot \text{tr}(\mathbf{G}\mathbf{A}) \cdot \mathbf{G}, \quad \mathbb{E}_{\mathcal{D}_{\text{target}}}[\mathbf{x}\mathbf{x}^\top \mathbf{A}\mathbf{x}\mathbf{x}^\top] - \mathbf{H}\mathbf{A}\mathbf{H} \succeq \beta \cdot \text{tr}(\mathbf{H}\mathbf{A}) \cdot \mathbf{H}.$$

Assumption 1 holds with $\alpha = 3$ and $\beta = 1$ given that $\mathcal{D}_{\text{source}}(\mathbf{x}) = \mathcal{N}(\mathbf{0}, \mathbf{G})$ and $\mathcal{D}_{\text{target}}(\mathbf{x}) = \mathcal{N}(\mathbf{0}, \mathbf{H})$. Moreover, Assumption 1A holds if both $\mathbf{H}^{-\frac{1}{2}} \cdot \mathcal{D}_{\text{source}}(\mathbf{x})$ and $\mathbf{G}^{-\frac{1}{2}} \cdot \mathcal{D}_{\text{target}}(\mathbf{x})$ have sub-Gaussian tails [Zou et al., 2021b]. For more exemplar distributions that satisfy Assumption 1, we refer the reader to Wu et al. [2021].

Assumption 2 (Noise condition). *Assume that there is a constant $\sigma^2 > 0$ such that*

$$\mathbb{E}_{\mathcal{D}_{\text{source}}}[(y - \langle \mathbf{w}^*, \mathbf{x} \rangle)^2 \mathbf{x}\mathbf{x}^\top] \preceq \sigma^2 \cdot \mathbf{G}, \quad \mathbb{E}_{\mathcal{D}_{\text{target}}}[(y - \langle \mathbf{w}^*, \mathbf{x} \rangle)^2 \mathbf{x}\mathbf{x}^\top] \preceq \sigma^2 \cdot \mathbf{H}.$$

Assumption 2 puts mild requirements on the conditional distribution of the response given input covariates for both source and target distribution. In particular, Assumption 2 is directly implied by the following Assumption 2' for a well-specified linear regression model under covariate shift.

²For the conciseness of presentation we focus on SGD with geometrically decaying stepsizes. With the provided techniques, our results can be easily extended to SGD with tail geometrically decaying stepsizes [Wu et al., 2021] as well.

Assumption 2' (Well-specified noise). Assume that for both source and target distributions, the response (conditional on input covariates) is given by

$$y = \mathbf{x}^\top \mathbf{w}^* + \epsilon, \text{ where } \epsilon \sim \mathcal{N}(0, \sigma^2) \text{ and } \epsilon \text{ is independent with } \mathbf{x}.$$

Additional Notation. Let $\mathbb{N}_+ := \{1, 2, \dots\}$. For an index set $\mathbb{K} \subset \mathbb{N}_+$, its complement is defined by $\mathbb{K}^c := \mathbb{N}_+ - \mathbb{K}$. Then for an index set $\mathbb{K} \subset \mathbb{N}_+$ and a scalar $a \geq 0$, we define

$$\mathbf{H}_{\mathbb{K}} := \sum_{i \in \mathbb{K}} \lambda_i \mathbf{v}_i \mathbf{v}_i^\top, \quad \mathbf{H}_{\mathbb{K}}^{-1} := \sum_{i \in \mathbb{K}} \frac{1}{\lambda_i} \mathbf{v}_i \mathbf{v}_i^\top, \quad a\mathbf{I}_{\mathbb{K}} + \mathbf{H}_{\mathbb{K}^c} := \sum_{i \in \mathbb{K}} a \mathbf{v}_i \mathbf{v}_i^\top + \sum_{i \notin \mathbb{K}} \lambda_i \mathbf{v}_i \mathbf{v}_i^\top,$$

where $(\lambda_i)_{i \geq 1}$ and $(\mathbf{v}_i)_{i \geq 1}$ are corresponding eigenvalues and eigenvectors of $\mathbf{H}$. One can verify that $\mathbf{H}_{\mathbb{K}}^{-1}$ is equivalent to the (pseudo) inverse of $\mathbf{H}_{\mathbb{K}}$. Similarly, we define $\mathbf{G}_{\mathbb{J}}$, $\mathbf{G}_{\mathbb{J}}^{-1}$ and $a\mathbf{I}_{\mathbb{J}} + \mathbf{G}_{\mathbb{J}^c}$ according to the eigenvalues and eigenvectors of $\mathbf{G}$.

3 Main Results

An Upper Bound. We begin with presenting an upper bound for the target domain excess risk achieved by the pretraining-finetuning method.

Theorem 3.1 (upper bound). Suppose that Assumptions 1A and 2 hold. Let $\mathbf{w}_{M+N}$ be the output of (SGD). Let $M_{\text{eff}} := M/\log(M)$, $N_{\text{eff}} := N/\log(N)$. Suppose that $\gamma_0, \gamma_M < \min\{1/(4\alpha \text{tr}(\mathbf{G})), 1/(4\alpha \text{tr}(\mathbf{H}))\}$. Then it holds that

$$\text{ExcessRisk}(\mathbf{w}_{M+N}) \leq \text{BiasError} + \text{VarError}.$$

Moreover, for any two index sets $\mathbb{J}, \mathbb{K} \subset \mathbb{N}_+$, it holds that

$$\begin{aligned} \text{VarError} &\lesssim \sigma^2 \cdot \left(\frac{D_{\text{eff}}^{\text{finetune}}}{M_{\text{eff}}} + \frac{D_{\text{eff}}}{N_{\text{eff}}} \right); \\ \text{BiasError} &\lesssim \left\| \prod_{t=M}^{M+N-1} (\mathbf{I} - \gamma_t \mathbf{H}) \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G}) (\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 \\ &\quad + \alpha \cdot \left\| \mathbf{w}_0 - \mathbf{w}^* \right\|_{\frac{\mathbf{I}_{\mathbb{J}}}{M_{\text{eff}}\gamma_0} + \mathbf{G}_{\mathbb{J}^c}}^2 \cdot \frac{D_{\text{eff}}^{\text{finetune}}}{M_{\text{eff}}} \\ &\quad + \alpha \cdot \left(\left\| \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G}) (\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\frac{\mathbf{I}_{\mathbb{K}}}{N_{\text{eff}}\gamma_M} + \mathbf{H}_{\mathbb{K}^c}}^2 + \left\| \mathbf{w}_0 - \mathbf{w}^* \right\|_{\frac{\mathbf{I}_{\mathbb{J}}}{M_{\text{eff}}\gamma_0} + \mathbf{G}_{\mathbb{J}^c}}^2 \right) \cdot \frac{D_{\text{eff}}}{N_{\text{eff}}}, \end{aligned}$$

where

$$\begin{aligned} D_{\text{eff}}^{\text{finetune}} &:= \text{tr} \left(\prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H})^2 \mathbf{H} \cdot (\mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}}^2 \gamma_0^2 \cdot \mathbf{G}_{\mathbb{J}^c}) \right), \\ D_{\text{eff}} &:= |\mathbb{K}| + N_{\text{eff}}^2 \gamma_M^2 \cdot \sum_{i \notin \mathbb{K}} \lambda_i^2. \end{aligned} \tag{1}$$

In particular, the upper bounds are optimized when

$$\mathbb{J} = \{j : \mu_j \geq 1/(\gamma_0 M_{\text{eff}})\}, \quad \mathbb{K} = \{k : \lambda_k \geq 1/(\gamma_M N_{\text{eff}})\}. \tag{2}$$

The upper bound in Theorem 3.1 contains a bias error stemming from the incorrect initialization $\mathbf{w}_0 \neq \mathbf{w}^*$, and a variance error caused by the additive label noise $y - \mathbf{x}^\top \mathbf{w}^* \neq 0$. In particular, M_{eff} and N_{eff} are the *effective number of source and target data*, respectively, due to the effect of the geometrically decaying stepsizes in (SGD). Moreover, D_{eff} can be regarded as the *effective dimension of supervised learning* [Wu et al., 2021] and $D_{\text{eff}}^{\text{finetune}}$ can be regarded as the *effective dimension of pretraining-finetuning*. Note that $D_{\text{eff}}^{\text{finetune}}$ is determined jointly by the spectrum of the source and target population covariance matrices as well as the stepsizes for pretraining and finetuning.

To better illustrate the spirit of Theorem 3.1, let us consider an example where $\|\mathbf{w}_0 - \mathbf{w}^*\|_2^2, \sigma^2 \lesssim 1$, $\gamma_0 \approx 1$, and $\text{tr}(\mathbf{G}) \approx \text{tr}(\mathbf{H}) \approx 1$ (so that the spectrum of $\mathbf{G}$ and $\mathbf{H}$ must decay fast), then the bound in Theorem 3.1 vanishes provided that

$$D_{\text{eff}} = o(N_{\text{eff}}), \quad D_{\text{eff}}^{\text{finetune}} = o(M_{\text{eff}}). \tag{3}$$

For the first condition in (3) to happen one needs

$$|\mathbb{K}| = o(N/\log N), \quad \gamma_M^2 \cdot \sum_{i \notin \mathbb{K}} \lambda_i^2 = o(\log N/N),$$

which can be satisfied when **(i)** the number of target data N is large and the finetuning stepsize $\gamma_M \approx 1$, or when **(ii)** N is small and γ_M is also small (which can depend on N). The second condition in (3) can happen under various situations, e.g., when **(i)** N is large and $\gamma_M \approx 1$, or when **(ii)** N, γ_M are small but the amount of source data M is large and that

$$\text{tr}(\mathbf{H}\mathbf{G}_{\mathbb{J}}^{-1}) = o(M/\log M), \quad \text{tr}(\mathbf{H}\mathbf{G}_{\mathbb{J}^c}) = o(\log M/M),$$

which will hold when $\mathbf{G}$ aligns well with $\mathbf{H}$ (as a sanity check these hold automatically when $\mathbf{G} = \mathbf{H}$ and M is large). To summarize, in case **(i)** the amount of target data is plentiful so that finetuning with large stepsize leads to generalization (which is essentially supervised learning); and in case **(ii)**, even though the target data is scarce, pretraining with abundant source data can still generalize given that the source and target population covariance matrices are well aligned.

A Lower Bound. The following theorem provides a nearly matching lower bound.

Theorem 3.2 (lower bound). *Suppose that Assumptions 1B and 2' hold. Let $\mathbf{w}_{M+N}$ be the output of (SGD). Let $M_{\text{eff}} := M/\log(M)$, $N_{\text{eff}} := N/\log(N)$, and suppose that $M_{\text{eff}}, N_{\text{eff}} \geq 10$. Suppose that $\gamma_0 < 1/\|\mathbf{G}\|_2$, $\gamma_M < 1/\|\mathbf{H}\|_2$. Then it holds that*

$$\text{ExcessRisk}(\mathbf{w}_{M+N}) = \text{BiasError} + \text{VarError}.$$

Moreover, for the index sets $\mathbb{K}$ and $\mathbb{J}$ defined in (2), it holds that

$$\begin{aligned} \text{VarError} &\gtrsim \sigma^2 \cdot \left(\frac{D_{\text{eff}}^{\text{finetune}}}{M_{\text{eff}}} + \frac{D_{\text{eff}}}{N_{\text{eff}}} \right); \\ \text{BiasError} &\gtrsim \left\| \prod_{t=M}^{M+N-1} (\mathbf{I} - \gamma_t \mathbf{H}) \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 \\ &\quad + \beta \cdot \|\mathbf{w}_0 - \mathbf{w}^*\|_{\mathbf{G}_{\mathbb{J}^c}}^2 \cdot \frac{D_{\text{eff}}^{\text{finetune}}}{M_{\text{eff}}} + \beta \cdot \left\| \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}_{\mathbb{K}^c}}^2 \cdot \frac{D_{\text{eff}}}{N_{\text{eff}}}, \end{aligned}$$

where D_{eff} and $D_{\text{eff}}^{\text{finetune}}$ are as defined in (1).

The lower bound in Theorem 3.2 suggests that the upper bound in Theorem 3.1 is tight upto constant factor in terms of variance error, and is also tight in terms of bias error except for the following additional parts in the respective places:

$$\left\| \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\frac{\mathbf{I}_{\mathbb{K}}}{N_{\text{eff}}\gamma_M}}^2, \quad \|\mathbf{w}_0 - \mathbf{w}^*\|_{\frac{\mathbf{I}_{\mathbb{J}}}{M_{\text{eff}}\gamma_0}}^2, \quad \|\mathbf{w}_0 - \mathbf{w}^*\|_{\frac{\mathbf{I}_{\mathbb{J}}}{M_{\text{eff}}\gamma_0} + \mathbf{G}_{\mathbb{J}^c}}^2.$$

In particular, the upper and lower bounds match ignoring constant factors provided that

$$\|\mathbf{w}_0 - \mathbf{w}^*\|_{\mathbf{G}}^2 \lesssim \sigma^2, \quad \left\| \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 \lesssim \sigma^2,$$

which hold in a statistically interesting regime where the signal-to-noise ratios, $\|\mathbf{w}_0 - \mathbf{w}^*\|_{\mathbf{G}}^2/\sigma^2$, $\|\mathbf{w}_0 - \mathbf{w}^*\|_{\mathbf{H}}^2/\sigma^2$, are bounded and $\mathbf{G}$ commutes with $\mathbf{H}$.

Implication for Pretraining. If target data is unavailable, Theorems 3.1 and 3.2 imply the following corollary for pretraining.

Corollary 3.3 (Learning with only source data). *Suppose that Assumptions 1A and 2 hold. Let $\mathbf{w}_{M+0}$ be the output of (SGD) with $M > 100$ source data and 0 target data. Suppose that $\gamma := \gamma_0 < 1/(4\alpha \text{tr}(\mathbf{H}))$. Then it holds that*

$$\text{ExcessRisk}(\mathbf{w}_{M+0}) \lesssim \left\| \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 + \left(\alpha \|\mathbf{w}_0 - \mathbf{w}^*\|_{\frac{\mathbf{I}_{\mathbb{J}}}{M_{\text{eff}}\gamma} + \mathbf{G}_{\mathbb{J}^c}}^2 + \sigma^2 \right) \frac{D_{\text{eff}}^{\text{pretrain}}}{M_{\text{eff}}},$$

where $D_{\text{eff}}^{\text{pretrain}} := \text{tr}(\mathbf{H}\mathbf{G}_{\mathbb{J}}^{-1}) + M_{\text{eff}}^2 \gamma^2 \cdot \text{tr}(\mathbf{H}\mathbf{G}_{\mathbb{J}^c})$ and $\mathbb{J} \subset \mathbb{N}_+$ can be any index set. If in addition Assumptions 1B and 2' hold, then for the index set $\mathbb{J}$ defined in (2), it holds that

$$\text{ExcessRisk}(\mathbf{w}_{M+0}) \gtrsim \left\| \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 + (\beta \|\mathbf{w}_0 - \mathbf{w}^*\|_{\mathbf{G}_{\mathbb{J}^c}}^2 + \sigma^2) \cdot \frac{D_{\text{eff}}^{\text{pretrain}}}{M_{\text{eff}}}.$$

Corollary 3.3 sharply characterizes the generalization of pretraining method, and is tight upto constant factors provided with a bounded signal-to-noise ratio, i.e., $\|\mathbf{w}_0 - \mathbf{w}^*\|_{\mathbf{G}}^2 \lesssim \sigma^2$. Corollary 3.3 can be interpreted in a similar way as Theorem 3.1. Moreover, these sharp bounds for pretraining and pretraining-finetuning enable us to study the effects of pretraining and finetuning thoroughly, which we will do in Section 4.

Implication for Supervised Learning. As a bonus, we can also apply Theorems 3.1 and 3.2 in the setting of supervised learning.

Corollary 3.4 (Learning with only target data). *Suppose that Assumptions 1A and 2 hold. Let $\mathbf{w}_{0+N}$ be the output of (SGD) with 0 source data and $N > 100$ target data. Suppose that $\gamma := \gamma_M < 1/(4\alpha(\text{tr}(\mathbf{G})))$. Then it holds that*

$$\text{ExcessRisk}(\mathbf{w}_{0+N}) \lesssim \left\| \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_t \mathbf{H})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 + \left(\alpha \left\| \mathbf{w}_0 - \mathbf{w}^* \right\|_{\frac{\mathbf{I}_{\mathbb{K}}}{N_{\text{eff}}\gamma} + \mathbf{H}_{\mathbb{K}^c}}^2 + \sigma^2 \right) \frac{D_{\text{eff}}}{N_{\text{eff}}},$$

where D_{eff} is as defined in (1) and $\mathbb{K} \subset \mathbb{N}_+$ can be any index set. If in addition Assumptions 1B and 2' hold, then for the index set $\mathbb{K}$ defined in (2), it holds that

$$\text{ExcessRisk}(\mathbf{w}_{0+N}) \gtrsim \left\| \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_t \mathbf{H})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 + (\beta \left\| \mathbf{w}_0 - \mathbf{w}^* \right\|_{\mathbf{H}_{\mathbb{K}^c}}^2 + \sigma^2) \cdot \frac{D_{\text{eff}}}{N_{\text{eff}}}.$$

We remark that the upper bound in Corollary 3.4 improves the related bound in Wu et al. [2021] by a logarithmic factor, and matches the lower bound upto constant factors given that $\left\| \mathbf{w}_0 - \mathbf{w}^* \right\|_{\mathbf{H}}^2 \lesssim \sigma^2$.

For a more detailed comparison between the bounds in Theorem 3.1, Corollaries 3.3 and 3.4, we refer the reader to Table 1 in Appendix A.

4 Discussions

With the established bounds, we are ready to discuss the power and limitation of pretraining and finetuning by comparing them to supervised learning.

The Power of Pretraining. For covariate shift problem, pretraining with infinite many source data can learn the true model. But when there are only a finite number of source data, it is unclear how the effect of pretraining compares to the effect of supervised learning (with finite many target data). Our next result quantitatively address this question by comparing Corollary 3.3 with Corollary 3.4.

Theorem 4.1 (Pretraining vs. supervised learning). *Suppose that Assumptions 1 and 2' hold. Let $\mathbf{w}_{0+N^{\text{sl}}}$ be the output of (SGD) with optimally tuned initial stepsize, 0 source data and $N^{\text{sl}} > 100$ target data. Let $\mathbf{w}_{M+0}$ be the output of (SGD) with optimally tuned initial stepsize, M source data and 0 target data. Let $M_{\text{eff}} := M/\log(M)$, $N_{\text{eff}}^{\text{sl}} := N^{\text{sl}}/\log(N^{\text{sl}})$. Suppose all SGD methods are initialized from $\mathbf{0}$. Then for every covariate shift problem instance $(\mathbf{w}^*, \mathbf{H}, \mathbf{G})$ such that $\mathbf{H}, \mathbf{G}$ commute, it holds that*

$$\text{ExcessRisk}(\mathbf{w}_{M+0}) \lesssim (1 + \alpha \left\| \mathbf{w}^* \right\|_{\mathbf{G}}^2 / \sigma^2) \cdot \text{ExcessRisk}(\mathbf{w}_{0+N^{\text{sl}}})$$

provided that

$$M_{\text{eff}} \geq (N_{\text{eff}}^{\text{sl}})^2 \cdot \frac{4 \left\| \mathbf{H}_{\mathbb{K}^*} \right\|_{\mathbf{G}}}{\alpha D_{\text{eff}}^{\text{sl}}},$$

where

$$\mathbb{K}^* := \{k : \lambda_k \geq 1/(N_{\text{eff}}^{\text{sl}} \gamma^{\text{sl}})\}, \quad D_{\text{eff}}^{\text{sl}} := |\mathbb{K}^*| + (N_{\text{eff}}^{\text{sl}} \gamma^{\text{sl}})^2 \sum_{i \notin \mathbb{K}^*} \lambda_i^2,$$

and $\gamma^{\text{sl}} < 1/\left\| \mathbf{H} \right\|_2$ refers to the optimal initial stepsize for supervised learning.

We now explain the implication of Theorem 4.1. First of all, it is of statistical interest to consider a signal-to-noise ratio bounded from above, i.e., $\left\| \mathbf{w}^* \right\|_{\mathbf{G}}^2 / \sigma^2 \lesssim 1$. Note that $D_{\text{eff}}^{\text{sl}} \geq 1$ when N^{sl} is large. Moreover, recall that $|\mathbb{K}^*| \leq D_{\text{eff}}^{\text{sl}} = o(N_{\text{eff}}^{\text{sl}})$ when supervised learning can achieve a vanishing excess risk. Finally, $\left\| \mathbf{H}_{\mathbb{K}^*} \right\|_{\mathbf{G}} := \left\| \mathbf{G}^{-\frac{1}{2}} \mathbf{H}_{\mathbb{K}^*} \mathbf{G}^{-\frac{1}{2}} \right\|_2 \lesssim 1$ can be satisfied if the top $|\mathbb{K}^*| = o(N_{\text{eff}}^{\text{sl}})$ eigenvalues subspace of $\mathbf{H}$ mostly falls into the top eigenvalues subspace of $\mathbf{G}$. Under these remarks, Theorem 4.1 suggests that: in the bounded signal-to-noise cases, pretraining with $O(N^2)$ source data is *no worse than* supervised learning with N target data (ignoring constant factors), for *every* covariate shift problem such that the *top* eigenvalues subspace of the target covariance matrix aligns well with that of the source covariance matrix.

The Power of Pretraining-Finetuning. We next discuss the effect of pretraining-finetuning by comparing Theorem 3.1 with Corollary 3.4.

Theorem 4.2 (Pretraining-finetuning vs. supervised learning). *Suppose that Assumptions 1 and 2' hold. Let $\mathbf{w}_{0+N^{\text{sl}}}$ be the output of (SGD) with optimally tuned initial stepsize, 0 source data and $N^{\text{sl}} > 100$ target data. Let $\mathbf{w}_{M+N}$ be the output of (SGD) with optimally tuned initial stepsize,*

M source data and N target data. Let $M_{\text{eff}} := M/\log(M)$, $N_{\text{eff}} := N/\log(N)$, $N_{\text{eff}}^{\text{sl}} := N^{\text{sl}}/\log(N^{\text{sl}})$. Suppose all SGD methods are initialized from $\mathbf{0}$. Then for every covariate shift problem instance $(\mathbf{w}^*, \mathbf{H}, \mathbf{G})$ such that $\mathbf{H}, \mathbf{G}$ commute, it holds that

$$\text{ExcessRisk}(\mathbf{w}_{M+N}) \lesssim (1 + \alpha \|\mathbf{w}^*\|_{\mathbf{G}}^2 / \sigma^2) \cdot \text{ExcessRisk}(\mathbf{w}_{0+N^{\text{sl}}})$$

provided that

$$M_{\text{eff}} \geq (N_{\text{eff}}^{\text{sl}})^2 \cdot \frac{4 \|\mathbf{H}_{\mathbb{K}^\dagger}\|_{\mathbf{G}}}{\alpha D_{\text{eff}}^{\text{sl}}},$$

where

$$\mathbb{K}^\dagger := \{k : N_{\text{eff}}^{\text{sl}} \log(N_{\text{eff}}^{\text{sl}}) \text{tr}(\mathbf{H}) / (N_{\text{eff}} D_{\text{eff}}^{\text{sl}}) > \lambda_k \geq 1 / (N_{\text{eff}}^{\text{sl}} \gamma^{\text{sl}})\} \subset \mathbb{K}^*,$$

and $\mathbb{K}^*$, $D_{\text{eff}}^{\text{sl}}$, γ^{sl} are as defined in Theorem 4.1.

Theorem 4.2 can be interpreted in a similar way as Theorem 4.1. The only difference is that Theorem 4.2 puts a *milder* condition regarding the alignment of $\mathbf{H}$ and $\mathbf{G}$ than Theorem 4.1. In particular, Theorem 4.2 only requires the “middle” eigenvalues subspace of $\mathbf{H}$ mostly falls into the top eigenvalues subspace of $\mathbf{G}$. Moreover, the index set $\mathbb{K}^\dagger$ shrinks (hence $\|\mathbf{H}_{\mathbb{K}^\dagger}\|_{\mathbf{G}}$ decreases) as the number of target data for finetuning increases. This indicates that finetuning can help save the amount of source data for pretraining.

The Limitation of Pretraining vs. the Power of Finetuning. The following example further demonstrates the limitation of pretraining and the power of finetuning.

Example 4.3 (Pretraining-finetuning vs. pretraining vs. supervised learning). *Let $\epsilon > 0$ be a sufficiently small constant. Consider a covariate shift problem instance given by*

$$\begin{aligned} \mathbf{w}^* &= (1, 1, 0, 0, \dots)^\top, \quad \sigma^2 = 1, \\ \mathbf{H} &= \text{diag}(1, \underbrace{\epsilon^{0.5}, \dots, \epsilon^{0.5}}_{2\epsilon^{-0.5} \text{ copies}}, 0, 0, \dots), \quad \mathbf{G} = \text{diag}(\epsilon^2, 1, 0, 0, \dots). \end{aligned}$$

One may verify that $\text{tr}(\mathbf{H}) \approx \text{tr}(\mathbf{G}) \approx 1$ and that $\|\mathbf{w}^*\|_{\mathbf{H}}^2 \approx \|\mathbf{w}^*\|_{\mathbf{G}}^2 \approx \sigma^2 \approx 1$. The following holds for the (SGD) output:

- **supervised learning:** for $\text{ExcessRisk}(\mathbf{w}_{0+N}) < \epsilon$, it is necessary to have that $N \gtrsim \epsilon^{-1.5}$;
- **pretraining:** for $\text{ExcessRisk}(\mathbf{w}_{M+0}) < \epsilon$, it is necessary to have that $M \gtrsim \epsilon^{-2}$;
- **pretrain-finetuning:** for $\text{ExcessRisk}(\mathbf{w}_{M+N}) < \epsilon$, it suffices to set $\gamma_0 \approx 1$, $\gamma_M \approx \epsilon$, and $M \approx \epsilon^{-1} \log(\epsilon^{-1})$, $N \approx \epsilon^{-1} \log^2(\epsilon^{-1})$.

It is clear that whenever target data are available, the optimally tuned pretraining-finetuning method is *always no worse than* the optimally tuned pretraining method, as one can simply set the finetuning stepsize to be small (or zero) so that the former reduces to the latter. Moreover, Example 4.3 shows a covariate shift problem instance such that pretraining-finetuning can save *polynomially* amount of data compared to pretraining (or supervised learning). This example demonstrates the limitation of pretraining and the benefits of finetuning. As a final remark for Example 4.3, direct computation implies that $\mathbb{K}^* = \{1, 2, \dots, 2\epsilon^{-0.5} + 1\}$ and $\mathbb{K}^\dagger = \{2, 3, \dots, 2\epsilon^{-0.5} + 1\}$, therefore $\|\mathbf{H}_{\mathbb{K}^*}\|_{\mathbf{G}} = \epsilon^{-2} \gg \|\mathbf{H}_{\mathbb{K}^\dagger}\|_{\mathbf{G}} = \epsilon^{0.5}$, so the implication of Example 4.3 is consistent with Theorems 4.1 and 4.2.

Numerical Simulations. We perform experiments on synthetic data to verify our theory. The code and data for our experiments can be found on Github³. Recall that the effectiveness of pretraining and finetuning depends on the alignment between source and target covariance matrices, therefore we design experiments where the source and target covariance matrices are aligned at different levels. In particular, we consider commutable matrices $\mathbf{G}$ and $\mathbf{H}$ with eigenvalues $\{\mu_i\}_{i \geq 1} = \{i^{-2}\}_{i \geq 1}$ and $\{\lambda_i\}_{i \geq 1} = \{i^{-1.5}\}_{i \geq 1}$, respectively. To simulate different alignments between $\mathbf{G}$ and $\mathbf{H}$, we first sort them so that both of their eigenvalues are in descending order, and then reverse the top- k eigenvalues

³<https://github.com/uclaml/pretrain-finetune-SGD>

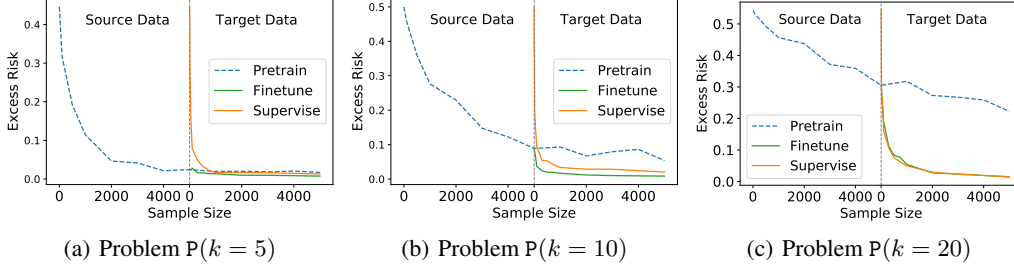


Figure 1: A generalization comparison between pretraining, pretraining-finetuning, and supervised learning. For each point in the curves, its x-axis represents the sample size and its y-axis represents the excess risk achieved by an algorithm with the corresponding amount of samples under the optimally tuned stepsizes. For the pretraining curves, the sample size refers to the amount of source data and the sample size appeared in the right half of the plots should be added by 5,000. The finetuning curves are generated from an initial model pretrained with 5,000 source data and its sample size refers to the amount of target data. The problem instances $P(k = 5)$, $P(k = 10)$, and $P(k = 20)$ are designed according to (4). The problem dimension is 200. The results are averaged over 20 independent repeats.

of $\mathbf{G}$. In mathematics, for a given $k > 0$, the problem instance $P(k) := (\mathbf{w}^*, \mathbf{H}, \mathbf{G}, \sigma^2)$ is designed as follows:

$$\mathbf{H} = \text{diag} \left(1, \frac{1}{2^{1.5}}, \dots, \frac{1}{k^{1.5}}, \frac{1}{(k+1)^{1.5}}, \dots \right), \quad \mathbf{G} = \text{diag} \left(\frac{1}{k^2}, \dots, \frac{1}{2^2}, 1, \frac{1}{(k+1)^2}, \dots \right),$$

$$\mathbf{w}^* = \left(\underbrace{1, 1, \dots, 1}_{k \text{ copies}}, 1/(k+1), 1/(k+2), \dots \right)^\top, \quad \sigma^2 = 1. \quad (4)$$

One can verify that $\text{tr}(\mathbf{H}) \approx \text{tr}(\mathbf{G}) \approx 1$ and that $\|\mathbf{w}^*\|_{\mathbf{H}}^2 \approx \|\mathbf{w}^*\|_{\mathbf{G}}^2 \approx \sigma^2 \approx 1$. Clearly, a larger k implies a worse alignment between $\mathbf{G}$ and $\mathbf{H}$. We then test three problem instances $P(k = 5)$, $P(k = 10)$, and $P(k = 20)$, and compare the excess risk achieved by pretraining, pretraining-finetuning, and supervised learning. The results are presented in Figure 1, which lead to the following informative observations:

- For problem $P(k = 5)$ where $\mathbf{H}$ and $\mathbf{G}$ are aligned very well, pretraining (without finetuning!) can already match the generalization performance of supervised learning. This verifies the power of pretraining for tackling transfer learning with mildly shifted covariate.
- For problem $P(k = 10)$ where $\mathbf{H}$ and $\mathbf{G}$ are moderately aligned, there is a significant gap between the risk of pretraining and that of supervised learning. Yet, the gap is closed when the pretrained model is finetuned with scarce target data. This demonstrates the limitation of pretraining and the power of finetuning for tackling transfer learning with moderate shifted covariate.
- For problem $P(k = 20)$ where $\mathbf{H}$ and $\mathbf{G}$ are poorly aligned, the risk of pretraining can hardly compete with that of supervised learning. Moreover, for finetuning to match the performance of supervised learning, it requires nearly the same amount of target data as that used by supervised learning. This reveals the limitation of pretraining and finetuning for tackling transfer learning with severely shifted covariate.

5 Concluding Remarks

We consider linear regression under covariate shift, and a SGD estimator that is firstly trained with source domain data and then finetuned with target domain data. We derive sharp upper and lower bounds for the estimator's target domain excess risk. Based on the derived bounds, we show that for a large class of covariate shift problems, pretraining with $O(N^2)$ source data can match the performance of supervised learning with N target data. Moreover, we show that finetuning with scarce target data can significantly reduce the amount of source data required by pretraining. Finally, when applied to supervised linear regression, our results improve the upper bound in [Wu et al., 2021] by a logarithmic factor, and close its gap with the lower bound (ignoring constant factors) when the signal-to-noise ratio is bounded.

Several future directions are worth discussing.

Model Shift. An immediate follow-up problem is to extend our results from the covariate shift setting to more general transfer learning settings, e.g., with both covariate shift and *model shift*, where the true parameter $\mathbf{w}^*$ could also be different for source and target distributions. Under model shift, the power of pretraining with source data is limited, and we expect that finetuning with target data becomes even more important.

Ridge Regression. For infinite-dimensional least-squares in the supervised learning context, instance-wisely tight bounds for both ridge regression and SGD have been established by Bartlett et al. [2020], Tsigler and Bartlett [2020], Zou et al. [2021b], Wu et al. [2021]. For infinite-dimensional least-squares under covariate shift, this paper presents nearly instance-wisely tight bounds for SGD. As ridge regression is popular in covariate shift literature (see Ma et al. [2022] and references herein), an interesting future direction is studying the instance-wisely tight bounds for ridge regression in the setting of infinite-dimensional least-squares under covariate shift — a tight bias analysis is of particular interest.

Unlabeled Data. In this work we assume the provided source and target data are both labeled. However in many practical scenarios, additional unlabeled source and target data are also available. In this case our results cannot be directly applied as it remains unclear how to utilize unlabeled data with SGD. An important future direction is to extend our framework to incorporate with unlabeled source and target data.

Acknowledgments and Disclosure of Funding

We would like to thank the anonymous reviewers and area chairs for their helpful comments. This work was done when DZ was a Ph.D. student at UCLA. DZ is partially supported by Bloomberg data science Ph.D. fellowship and his startup funding in the institute of data science, the University of Hong Kong. JW and VB are supported by the Defense Advanced Research Projects Agency (DARPA) under Contract No. HR00112190130. QG is partially supported by the National Science Foundation award IIS-1906169 and IIS-2008981. SK acknowledges funding from the Office of Naval Research under award N00014-22-1-2377 and the National Science Foundation Grant under award CCF-1703574. The views and conclusions contained in this paper are those of the authors and should not be interpreted as representing any funding agencies.

References

- Martin Arjovsky, Léon Bottou, Ishaan Gulrajani, and David Lopez-Paz. Invariant risk minimization. *arXiv preprint arXiv:1907.02893*, 2019.
- Francis Bach and Eric Moulines. Non-strongly-convex smooth stochastic approximation with convergence rate $o(1/n)$. *Advances in neural information processing systems*, 26:773–781, 2013.
- Raghu Raj Bahadur. Rates of convergence of estimates and test statistics. *The Annals of Mathematical Statistics*, 38(2):303–324, 1967.
- Raghu Raj Bahadur. *Some limit theorems in statistics*. SIAM, 1971.
- Peter L Bartlett, Philip M Long, Gábor Lugosi, and Alexander Tsigler. Benign overfitting in linear regression. *Proceedings of the National Academy of Sciences*, 2020.
- Shai Ben-David, John Blitzer, Koby Crammer, Alex Kulesza, Fernando Pereira, and Jennifer Wortman Vaughan. A theory of learning from different domains. *Machine learning*, 79(1):151–175, 2010.
- Corinna Cortes and Mehryar Mohri. Domain adaptation and sample bias correction theory and algorithm for regression. *Theoretical Computer Science*, 519:103–126, 2014.
- Corinna Cortes, Yishay Mansour, and Mehryar Mohri. Learning bounds for importance weighting. *Advances in neural information processing systems*, 23, 2010.
- Corinna Cortes, Mehryar Mohri, and Andrés Muñoz Medina. Adaptation based on generalized discrepancy. *The Journal of Machine Learning Research*, 20(1):1–30, 2019.

- Shai Ben David, Tyler Lu, Teresa Luu, and Dávid Pál. Impossibility theorems for domain adaptation. In *Proceedings of the Thirteenth International Conference on Artificial Intelligence and Statistics*, pages 129–136. JMLR Workshop and Conference Proceedings, 2010.
- Aymeric Dieuleveut, Nicolas Flammarion, and Francis Bach. Harder, better, faster, stronger convergence rates for least-squares regression. *The Journal of Machine Learning Research*, 18(1): 3520–3570, 2017.
- Rong Ge, Sham M Kakade, Rahul Kidambi, and Praneeth Netrapalli. The step decay schedule: A near optimal, geometrically decaying learning rate procedure for least squares. *arXiv preprint arXiv:1904.12838*, 2019.
- Pascal Germain, Amaury Habrard, François Laviolette, and Emilie Morvant. A pac-bayesian approach for domain adaptation with specialization to linear classifiers. In *International conference on machine learning*, pages 738–746. PMLR, 2013.
- Steve Hanneke and Samory Kpotufe. On the value of target data in transfer learning. *Advances in Neural Information Processing Systems*, 32, 2019.
- Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. *corr abs/1512.03385* (2015), 2015.
- Prateek Jain, Sham M Kakade, Rahul Kidambi, Praneeth Netrapalli, Venkata Krishna Pillutla, and Aaron Sidford. A markov chain theory approach to characterizing the minimax optimality of stochastic gradient descent (for least squares). *arXiv preprint arXiv:1710.09430*, 2017a.
- Prateek Jain, Praneeth Netrapalli, Sham M Kakade, Rahul Kidambi, and Aaron Sidford. Parallelizing stochastic gradient descent for least squares regression: mini-batching, averaging, and model misspecification. *The Journal of Machine Learning Research*, 18(1):8258–8299, 2017b.
- Samory Kpotufe and Guillaume Martinet. Marginal singularity, and the benefits of labels in covariate-shift. In *Conference On Learning Theory*, pages 1882–1886. PMLR, 2018.
- Qi Lei, Wei Hu, and Jason Lee. Near-optimal linear regression under distribution shift. In *International Conference on Machine Learning*, pages 6164–6174. PMLR, 2021.
- Cong Ma, Reese Pathak, and Martin J Wainwright. Optimally tackling covariate shift in rkhs-based nonparametric regression. *arXiv preprint arXiv:2205.02986*, 2022.
- Yishay Mansour, Mehryar Mohri, and Afshin Rostamizadeh. Domain adaptation: Learning bounds and algorithms. *arXiv preprint arXiv:0902.3430*, 2009.
- Mehryar Mohri and Andres Muñoz Medina. New analysis and algorithm for learning with drifting distributions. In *International Conference on Algorithmic Learning Theory*, pages 124–138. Springer, 2012.
- Sinno Jialin Pan and Qiang Yang. A survey on transfer learning. *IEEE Transactions on knowledge and data engineering*, 22(10):1345–1359, 2009.
- Reese Pathak, Cong Ma, and Martin J Wainwright. A new similarity measure for covariate shift with applications to nonparametric regression. *arXiv preprint arXiv:2202.02837*, 2022.
- Bernhard Schölkopf, Alexander J Smola, Francis Bach, et al. *Learning with kernels: support vector machines, regularization, optimization, and beyond*. MIT press, 2002.
- Hidetoshi Shimodaira. Improving predictive inference under covariate shift by weighting the log-likelihood function. *Journal of statistical planning and inference*, 90(2):227–244, 2000.
- Masashi Sugiyama and Motoaki Kawanabe. *Machine learning in non-stationary environments: Introduction to covariate shift adaptation*. MIT press, 2012.
- Masashi Sugiyama, Taiji Suzuki, and Takafumi Kanamori. *Density ratio estimation in machine learning*. Cambridge University Press, 2012.

- Alexander Tsigler and Peter L Bartlett. Benign overfitting in ridge regression. *arXiv preprint arXiv:2009.14286*, 2020.
- Aditya Varre, Loucas Pillaud-Vivien, and Nicolas Flammarion. Last iterate convergence of sgd for least-squares in the interpolation regime. *arXiv preprint arXiv:2102.03183*, 2021.
- Jingfeng Wu, Difan Zou, Vladimir Braverman, Quanquan Gu, and Sham M Kakade. Last iterate risk bounds of sgd with decaying stepsize for overparameterized linear regression. *arXiv preprint arXiv:2110.06198*, 2021.
- Yifan Wu, Ezra Winston, Divyansh Kaushik, and Zachary Lipton. Domain adaptation with asymmetrically-relaxed distribution alignment. In *International Conference on Machine Learning*, pages 6872–6881. PMLR, 2019.
- Jason Yosinski, Jeff Clune, Yoshua Bengio, and Hod Lipson. How transferable are features in deep neural networks? *Advances in neural information processing systems*, 27, 2014.
- Difan Zou, Jingfeng Wu, Vladimir Braverman, Quanquan Gu, Dean P Foster, and Sham Kakade. The benefits of implicit regularization from sgd in least squares problems. *Advances in Neural Information Processing Systems*, 34:5456–5468, 2021a.
- Difan Zou, Jingfeng Wu, Vladimir Braverman, Quanquan Gu, and Sham Kakade. Benign overfitting of constant-stepsize sgd for linear regression. In *Conference on Learning Theory*, pages 4633–4635. PMLR, 2021b.

Checklist

1. For all authors...
 - (a) Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope? [\[Yes\]](#)
 - (b) Did you describe the limitations of your work? [\[Yes\]](#)
 - (c) Did you discuss any potential negative societal impacts of your work? [\[N/A\]](#)
 - (d) Have you read the ethics review guidelines and ensured that your paper conforms to them? [\[Yes\]](#)
2. If you are including theoretical results...
 - (a) Did you state the full set of assumptions of all theoretical results? [\[Yes\]](#)
 - (b) Did you include complete proofs of all theoretical results? [\[Yes\]](#)
3. If you ran experiments...
 - (a) Did you include the code, data, and instructions needed to reproduce the main experimental results (either in the supplemental material or as a URL)? [\[Yes\]](#)
 - (b) Did you specify all the training details (e.g., data splits, hyperparameters, how they were chosen)? [\[Yes\]](#)
 - (c) Did you report error bars (e.g., with respect to the random seed after running experiments multiple times)? [\[No\]](#) The reported mean behaviors in our experiments are sufficient to support our theoretical results
 - (d) Did you include the total amount of compute and the type of resources used (e.g., type of GPUs, internal cluster, or cloud provider)? [\[No\]](#) Resources are not a concern for our purpose of numerically verifying our theoretical results
4. If you are using existing assets (e.g., code, data, models) or curating/releasing new assets...
 - (a) If your work uses existing assets, did you cite the creators? [\[N/A\]](#)
 - (b) Did you mention the license of the assets? [\[N/A\]](#)
 - (c) Did you include any new assets either in the supplemental material or as a URL? [\[N/A\]](#)
 - (d) Did you discuss whether and how consent was obtained from people whose data you’re using/curating? [\[N/A\]](#)

- (e) Did you discuss whether the data you are using/curating contains personally identifiable information or offensive content? [N/A]
- 5. If you used crowdsourcing or conducted research with human subjects...
 - (a) Did you include the full text of instructions given to participants and screenshots, if applicable? [N/A]
 - (b) Did you describe any potential participant risks, with links to Institutional Review Board (IRB) approvals, if applicable? [N/A]
 - (c) Did you include the estimated hourly wage paid to participants and the total amount spent on participant compensation? [N/A]

A A Comparison of Pretraining-Finetuning, Pretraining and Supervised Learning

In Table 1, we make a detailed comparison of the bounds for (1) pretraining-finetuning with M source data and N target data, (2) pretraining with M source data, and (3) supervised learning with N target data. The presented bounds are from Theorem 3.1, Corollaries 3.3 and 3.4. For simplicity, we assume that all SGD iterates are initialized from $\mathbf{w}_0 = 0$, and that the signal-to-noise ratios are bounded from above.

Table 1: A comparison of pretraining-finetuning, pretraining and supervised learning.

	Pretraining-Finetuning	Pretraining	Supervised Learning
initial stepsizes	γ_0, γ_M	γ_0	$\gamma_0 (= \gamma_M)$
number of data	$M + N$	$M + 0$	$0 + N$
effective number of source data (M_{eff})	$\frac{M}{\log(M)}$	$\frac{M}{\log(M)}$	0
effective number of target data (N_{eff})	$\frac{N}{\log(N)}$	0	$\frac{N}{\log(N)}$
source effective dimension ($D_{\text{eff}}^{\text{source}}$)	$\text{tr} \left(\prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H})^2 \mathbf{H} \cdot (\mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}}^2 \gamma_0^2 \mathbf{G}_{\mathbb{J}^c}) \right)$	$\text{tr}(\mathbf{H} \mathbf{G}_{\mathbb{J}}^{-1}) + M_{\text{eff}}^2 \gamma_0^2 \text{tr}(\mathbf{H} \mathbf{G}_{\mathbb{J}^c})$	0
target effective dimension ($D_{\text{eff}}^{\text{target}}$)	$ \mathbb{K} + N_{\text{eff}}^2 \gamma_M^2 \sum_{i \notin \mathbb{K}} \lambda_i^2$	0	$ \mathbb{K} + N_{\text{eff}}^2 \gamma_0^2 \sum_{i \notin \mathbb{K}} \lambda_i^2$
learnable indexes	$\mathbb{J} = \{j : \mu_j \geq 1/(\gamma_0 M_{\text{eff}})\}, \quad \mathbb{K} = \{k : \lambda_k \geq 1/(\gamma_M N_{\text{eff}})\}$		
Signal to noise ratio (SNR)	$\frac{\ \mathbf{w}^*\ _{\mathbf{G}}^2 + \ \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G}) \mathbf{w}^*\ _{\mathbf{H}}^2}{\sigma^2}$	$\frac{\ \mathbf{w}^*\ _{\mathbf{G}}^2}{\sigma^2}$	$\frac{\ \mathbf{w}^*\ _{\mathbf{H}}^2}{\sigma^2}$
effective bias error (bias_{eff})	$\left\ \prod_{t=M}^{M+N-1} (\mathbf{I} - \gamma_t \mathbf{H}) \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G}) \mathbf{w}^* \right\ _{\mathbf{H}}^2$	$\left\ \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_t \mathbf{G}) \mathbf{w}^* \right\ _{\mathbf{H}}^2$	$\left\ \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_t \mathbf{H}) \mathbf{w}^* \right\ _{\mathbf{H}}^2$
unified risk bound	$\text{bias}_{\text{eff}} + (1 + \sigma^2) \cdot \text{SNR} \cdot \left(\frac{D_{\text{eff}}^{\text{source}}}{M_{\text{eff}}} + \frac{D_{\text{eff}}^{\text{target}}}{N_{\text{eff}}} \right)$		

B Preliminaries

Notations. Define the following operators on symmetric matrices:

$$\mathcal{I} := \mathbf{I} \otimes \mathbf{I},$$

$$\mathcal{M}_{\mathbf{G}} := \mathbb{E}_{\text{source}}[\mathbf{x} \otimes \mathbf{x} \otimes \mathbf{x} \otimes \mathbf{x}], \quad \mathcal{M}_{\mathbf{H}} := \mathbb{E}_{\text{target}}[\mathbf{x} \otimes \mathbf{x} \otimes \mathbf{x} \otimes \mathbf{x}],$$

$$\begin{aligned}
\widetilde{\mathcal{M}}_{\mathbf{G}} &:= \mathbf{G} \otimes \mathbf{G}, & \widetilde{\mathcal{M}}_{\mathbf{H}} &:= \mathbf{H} \otimes \mathbf{H}, \\
\mathcal{T}_{\mathbf{G}}(\gamma) &:= \mathbf{G} \otimes \mathbf{I} + \mathbf{I} \otimes \mathbf{G} - \gamma \mathcal{M}_{\mathbf{G}}, & \mathcal{T}_{\mathbf{H}}(\gamma) &:= \mathbf{H} \otimes \mathbf{I} + \mathbf{I} \otimes \mathbf{H} - \gamma \mathcal{M}_{\mathbf{H}}, \\
\widetilde{\mathcal{T}}_{\mathbf{G}}(\gamma) &:= \mathbf{G} \otimes \mathbf{I} + \mathbf{I} \otimes \mathbf{G} - \gamma \mathbf{G} \otimes \mathbf{G}, & \widetilde{\mathcal{T}}_{\mathbf{H}}(\gamma) &:= \mathbf{H} \otimes \mathbf{I} + \mathbf{I} \otimes \mathbf{H} - \gamma \mathbf{H} \otimes \mathbf{H}
\end{aligned}$$

For the linear operators we have the following technical lemma from Zou et al. [2021b].

Lemma B.1 (Lemma B.1, Zou et al. [2021b]). *An operator $\mathcal{O}$ defined on symmetric matrices is called PSD mapping, if $\mathbf{A} \succeq 0$ implies $\mathcal{O} \circ \mathbf{A} \succeq 0$. Then we have*

1. $\mathcal{M}_{\mathbf{G}}, \mathcal{M}_{\mathbf{H}}, \widetilde{\mathcal{M}}_{\mathbf{G}}$ and $\widetilde{\mathcal{M}}_{\mathbf{H}}$ are all PSD mappings.
2. $\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma), \mathcal{I} - \gamma \mathcal{T}_{\mathbf{H}}(\gamma), \mathcal{I} - \gamma \widetilde{\mathcal{T}}_{\mathbf{G}}(\gamma)$ and $\mathcal{I} - \gamma \widetilde{\mathcal{T}}_{\mathbf{H}}(\gamma)$ are all PSD mappings.
3. $\mathcal{M}_{\mathbf{G}} - \widetilde{\mathcal{M}}_{\mathbf{G}}, \mathcal{M}_{\mathbf{H}} - \widetilde{\mathcal{M}}_{\mathbf{H}}, \widetilde{\mathcal{T}}_{\mathbf{G}} - \mathcal{T}_{\mathbf{G}}$ and $\widetilde{\mathcal{T}}_{\mathbf{H}} - \mathcal{T}_{\mathbf{H}}$ are all PSD mappings.
4. If $0 < \gamma < 1/\|\mathbf{G}\|_2$, then $\widetilde{\mathcal{T}}_{\mathbf{G}}^{-1}$ exists, and is a PSD mapping. Similarly, if $0 < \gamma < 1/\|\mathbf{H}\|_2$, then $\widetilde{\mathcal{T}}_{\mathbf{H}}^{-1}$ exists, and is a PSD mapping.
5. If $0 < \gamma < 1/(\alpha \text{tr}(\mathbf{G}))$, then $\mathcal{T}_{\mathbf{G}}^{-1} \circ \mathbf{A}$ exists for PSD matrix $\mathbf{A}$, and $\mathcal{T}_{\mathbf{G}}^{-1}$ is a PSD mapping. Similarly, if $0 < \gamma < 1/(\alpha \text{tr}(\mathbf{H}))$, then $\mathcal{T}_{\mathbf{H}}^{-1} \circ \mathbf{A}$ exists for PSD matrix $\mathbf{A}$, and $\mathcal{T}_{\mathbf{H}}^{-1}$ is a PSD mapping.
6. For every $\gamma > 0$ and every PSD matrices $\mathbf{A}$ and $\mathbf{B}$, we have

$$\begin{aligned}
\langle \mathbf{A}, (\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma)) \circ \mathbf{B} \rangle &= \langle (\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma)) \circ \mathbf{A}, \mathbf{B} \rangle, \\
\langle \mathbf{A}, (\mathcal{I} - \gamma \mathcal{T}_{\mathbf{H}}(\gamma)) \circ \mathbf{B} \rangle &= \langle (\mathcal{I} - \gamma \mathcal{T}_{\mathbf{H}}(\gamma)) \circ \mathbf{A}, \mathbf{B} \rangle.
\end{aligned}$$

Proof. Proof to the first five claims can be found in Lemma B.1 in Zou et al. [2021b]. The last claim is by definition. $\square$

Define

$$\boldsymbol{\Sigma}_{\mathbf{G}} := \mathbb{E}_{\text{source}}[(y - \langle \mathbf{w}^*, \mathbf{x} \rangle)^2 \mathbf{x} \mathbf{x}^\top], \quad \boldsymbol{\Sigma}_{\mathbf{H}} := \mathbb{E}_{\text{source}}[(y - \langle \mathbf{w}^*, \mathbf{x} \rangle)^2 \mathbf{x} \mathbf{x}^\top]$$

Then for the SGD iterates, we can consider their associated bias iterates and variance iterates:

$$\begin{cases} \mathbf{B}_0 = (\mathbf{w}_0 - \mathbf{w}^*)(\mathbf{w}_0 - \mathbf{w}^*)^\top; \\ \mathbf{B}_t = (\mathcal{I} - \gamma_{t-1} \mathcal{T}_{\mathbf{G}}(\gamma_{t-1})) \circ \mathbf{B}_{t-1}, & t = 1, \dots, M; \\ \mathbf{B}_{M+t} = (\mathcal{I} - \gamma_{M+t-1} \mathcal{T}_{\mathbf{H}}(\gamma_{M+t-1})) \circ \mathbf{B}_{M+t-1}, & t = 1, \dots, N; \end{cases} \quad (5)$$

$$\begin{cases} \mathbf{C}_0 = \mathbf{0}; \\ \mathbf{C}_t = (\mathcal{I} - \gamma_{t-1} \mathcal{T}_{\mathbf{G}}(\gamma_{t-1})) \circ \mathbf{C}_{t-1} + \gamma_t^2 \boldsymbol{\Sigma}_{\mathbf{G}}, & t = 1, \dots, M; \\ \mathbf{C}_{M+t} = (\mathcal{I} - \gamma_{M+t-1} \mathcal{T}_{\mathbf{H}}(\gamma_{M+t-1})) \circ \mathbf{C}_{M+t-1} + \gamma_{M+t-1}^2 \boldsymbol{\Sigma}_{\mathbf{H}}, & t = 1, \dots, N. \end{cases} \quad (6)$$

Lemma B.2 (Bias-variance decomposition). *Suppose that Assumption 2 holds. Then we have*

$$\mathbb{E}[\text{ExcessRisk}(\mathbf{w}_{M+N})] \leq \langle \mathbf{H}, \mathbf{B}_{M+N} \rangle + \langle \mathbf{H}, \mathbf{C}_{M+N} \rangle.$$

Proof. This follows from Lemma 2 in Wu et al. [2021]. $\square$

Lemma B.3 (Bias-variance decomposition, lower bound). *Suppose that Assumption 2' holds. Then we have*

$$\mathbb{E}[\text{ExcessRisk}(\mathbf{w}_{M+N})] = \frac{1}{2} \langle \mathbf{H}, \mathbf{B}_{M+N} \rangle + \frac{1}{2} \langle \mathbf{H}, \mathbf{C}_{M+N} \rangle.$$

Proof. This follows from Lemma 3 in Wu et al. [2021]. $\square$

C Variance Error Analysis

C.1 Upper Bounds

The following Assumption 1' is implied by Assumption 1A by setting $R^2 = \max\{\alpha \operatorname{tr}(\mathbf{H}), \alpha \operatorname{tr}(\mathbf{G})\}$. In this part we will work with the weaker Assumption 1'.

Assumption 1' (Fourth moment condition, relaxed version). *There exists a constant $R > 0$ such that*

$$\mathbb{E}_{\text{source}}[\mathbf{x}\mathbf{x}^\top \mathbf{x}\mathbf{x}^\top] \preceq R^2 \mathbf{G}, \quad \mathbb{E}_{\text{target}}[\mathbf{x}\mathbf{x}^\top \mathbf{x}\mathbf{x}^\top] \preceq R^2 \mathbf{H}.$$

Lemma C.1 (Crude bound on the variance iterates). *Suppose that Assumptions 1' and 2 hold. Suppose that $\max\{\gamma_0, \gamma_M\} \leq \gamma < 1/R^2$. Then it holds that*

$$\mathbf{C}_t \leq \frac{\gamma \sigma^2}{1 - \gamma R^2} \mathbf{I}, \text{ for every } t = 0, 1, \dots, M + N.$$

Proof. The proof idea has appeared in Jain et al. [2017a], Ge et al. [2019], Wu et al. [2021]. We prove the lemma by induction. For $t = 0$, it is clear that $\mathbf{C}_0 = \mathbf{0} \preceq \frac{\gamma \sigma^2}{1 - \gamma R^2} \mathbf{I}$. Now suppose that $\mathbf{C}_{t-1} \leq \frac{\gamma \sigma^2}{1 - \gamma R^2} \mathbf{I}$, and consider $\mathbf{C}_t$ according to (6). If $t \leq M$, then according to (6) we have

$$\begin{aligned} \mathbf{C}_t &= (\mathcal{I} - \gamma_{t-1} \mathcal{T}_{\mathbf{G}}(\gamma_{t-1})) \circ \mathbf{C}_{t-1} + \gamma_{t-1}^2 \Sigma_{\mathbf{G}} \\ &\preceq \frac{\gamma \sigma^2}{1 - \gamma R^2} \cdot (\mathcal{I} - \gamma_{t-1} \mathcal{T}_{\mathbf{G}}(\gamma_{t-1})) \circ \mathbf{I} + \gamma_{t-1}^2 \sigma^2 \cdot \mathbf{G} \\ &\preceq \frac{\gamma \sigma^2}{1 - \gamma R^2} \cdot (\mathbf{I} - 2\gamma_{t-1} \mathbf{G} + \gamma_{t-1}^2 \cdot R^2 \cdot \mathbf{G}) + \gamma_{t-1}^2 \sigma^2 \cdot \mathbf{G} \\ &= \frac{\gamma \sigma^2}{1 - \gamma R^2} \cdot \mathbf{I} - (2\gamma_{t-1} \gamma - \gamma_{t-1}^2) \cdot \frac{\sigma^2}{1 - \gamma R^2} \cdot \mathbf{G} \\ &\preceq \frac{\gamma \sigma^2}{1 - \gamma R^2} \cdot \mathbf{I}. \end{aligned} \tag{7}$$

If $t > M$, similarly according to (6) we have

$$\begin{aligned} \mathbf{C}_t &= (\mathcal{I} - \gamma_{t-1} \mathcal{T}_{\mathbf{H}}(\gamma_{t-1})) \circ \mathbf{C}_{t-1} + \gamma_{t-1}^2 \Sigma_{\mathbf{H}} \\ &\preceq \frac{\gamma \sigma^2}{1 - \gamma R^2} \cdot (\mathcal{I} - \gamma_{t-1} \mathcal{T}_{\mathbf{H}}(\gamma_{t-1})) \circ \mathbf{I} + \gamma_{t-1}^2 \sigma^2 \cdot \mathbf{H} \\ &\preceq \frac{\gamma \sigma^2}{1 - \gamma R^2} \cdot (\mathbf{I} - 2\gamma_{t-1} \mathbf{H} + \gamma_{t-1}^2 \cdot R^2 \cdot \mathbf{H}) + \gamma_{t-1}^2 \sigma^2 \cdot \mathbf{H} \\ &= \frac{\gamma \sigma^2}{1 - \gamma R^2} \cdot \mathbf{I} - (2\gamma_{t-1} \gamma - \gamma_{t-1}^2) \cdot \frac{\sigma^2}{1 - \gamma R^2} \mathbf{H} \\ &\preceq \frac{\gamma \sigma^2}{1 - \gamma R^2} \cdot \mathbf{I}. \end{aligned} \tag{8}$$

Putting everything together we complete the induction. $\square$

Lemma C.2 (Upper bounds on the variance iterates). *Suppose that Assumptions 1' and 2 hold. Suppose that $\max\{\gamma_0, \gamma_M\} \leq \gamma < 1/R^2$. Let $M_{\text{eff}} := M/\log(M)$, $N_{\text{eff}} := N/\log(N)$.*

- For every index set $\mathbb{J} \subset \mathbb{N}_+$, it holds that

$$\mathbf{C}_M \preceq \frac{8\sigma^2}{1 - \gamma R^2} \cdot \left(\frac{1}{M_{\text{eff}}} \cdot \mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}} \gamma_0^2 \cdot \mathbf{G}_{\mathbb{J}^c} \right).$$

- For every index set $\mathbb{K} \subset \mathbb{N}_+$, it holds that

$$\sum_{t=0}^{N-1} \gamma_{M+t}^2 \prod_{i=t+1}^{N-1} (\mathbf{I} - \gamma_{M+i} \mathbf{H})^2 \mathbf{H} \preceq 8 \cdot \left(\frac{1}{N_{\text{eff}}} \mathbf{H}_{\mathbb{K}}^{-1} + N_{\text{eff}} \gamma_M^2 \cdot \mathbf{H}_{\mathbb{K}^c} \right).$$

Proof. These are from the proof of Theorem 5 in Wu et al. [2021]. $\square$

Theorem C.1 (Variance error upper bound). *Suppose that Assumptions 1' and 2 hold. Suppose that $\max\{\gamma_0, \gamma_M\} \leq \gamma < 1/R^2$. Let $M_{\text{eff}} := M/\log(M)$, $N_{\text{eff}} := N/\log(N)$. Then it holds that*

$$\langle \mathbf{H}, \mathbf{C}_{M+N} \rangle \leq \frac{8\sigma^2}{1-\gamma R^2} \cdot \left(\frac{D_{\text{eff}}^{\text{finetune}}}{M_{\text{eff}}} + \frac{D_{\text{eff}}}{N_{\text{eff}}} \right),$$

where

$$D_{\text{eff}} := \text{tr}(\mathbf{H}\mathbf{H}_{\mathbb{K}}^{-1}) + N_{\text{eff}}^2 \gamma_M^2 \cdot \text{tr}(\mathbf{H}\mathbf{H}_{\mathbb{K}^c}),$$

$$D_{\text{eff}}^{\text{finetune}} := \text{tr} \left(\prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H})^2 \mathbf{H} \cdot (\mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}}^2 \gamma_0^2 \cdot \mathbf{G}_{\mathbb{J}^c}) \right),$$

and $\mathbb{K}, \mathbb{J}$ can be arbitrary index sets.

Proof of Theorem C.1. The core idea is to relate $\mathbf{C}_{M+N}$ to $\mathbf{C}_M$ via (6). For every $t = 0, \dots, N-1$, according to (6) we have

$$\begin{aligned} \mathbf{C}_{M+t+1} &= (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{C}_{M+t} + \gamma_{M+t}^2 \Sigma_{\mathbf{H}} \\ &\preceq (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{C}_{M+t} + \gamma_{M+t}^2 \cdot \mathcal{M}_{\mathbf{H}} \circ \mathbf{C}_{M+t} + \gamma_{M+t}^2 \sigma^2 \cdot \mathbf{H} \\ &\preceq (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{C}_{M+t} + \frac{\gamma_{M+t}^2 R^2 \cdot \gamma \sigma^2}{1-\gamma R^2} \cdot \mathbf{H} + \gamma_{M+t}^2 \sigma^2 \cdot \mathbf{H} \text{ (by Lemma C.1)} \\ &= (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{C}_{M+t} + \frac{\gamma_{M+t}^2 \sigma^2}{1-\gamma R^2} \cdot \mathbf{H}. \end{aligned}$$

Solving the above recursion from $t = 0$ to $t = N-1$ we obtain

$$\begin{aligned} \mathbf{C}_{M+N} &\preceq \prod_{t=0}^{N-1} (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{C}_M \\ &\quad + \frac{\sigma^2}{1-\gamma R^2} \cdot \sum_{t=0}^{N-1} \gamma_{M+t}^2 \prod_{i=t+1}^{N-1} (\mathcal{I} - \gamma_{M+i} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+i})) \circ \mathbf{H} \\ &= \prod_{t=0}^{N-1} (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{C}_M + \frac{\sigma^2}{1-\gamma R^2} \cdot \sum_{t=0}^{N-1} \gamma_{M+t}^2 \prod_{i=t+1}^{N-1} (\mathbf{I} - \gamma_{M+i} \mathbf{H})^2 \mathbf{H}. \end{aligned} \tag{9}$$

Therefore the variance error is

$$\begin{aligned} &\langle \mathbf{H}, \mathbf{C}_{M+N} \rangle \\ &\leq \left\langle \mathbf{H}, \prod_{t=0}^{N-1} (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{C}_M \right\rangle + \frac{\sigma^2}{1-\gamma R^2} \left\langle \mathbf{H}, \sum_{t=0}^{N-1} \gamma_{M+t}^2 \prod_{i=t+1}^{N-1} (\mathbf{I} - \gamma_{M+i} \mathbf{H})^2 \mathbf{H} \right\rangle \\ &= \left\langle \prod_{t=0}^{N-1} (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{H}, \mathbf{C}_M \right\rangle + \frac{\sigma^2}{1-\gamma R^2} \left\langle \mathbf{H}, \sum_{t=0}^{N-1} \gamma_{M+t}^2 \prod_{i=t+1}^{N-1} (\mathbf{I} - \gamma_{M+i} \mathbf{H})^2 \mathbf{H} \right\rangle \\ &= \left\langle \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H})^2 \mathbf{H}, \mathbf{C}_M \right\rangle + \frac{\sigma^2}{1-\gamma R^2} \left\langle \mathbf{H}, \sum_{t=0}^{N-1} \gamma_{M+t}^2 \prod_{i=t+1}^{N-1} (\mathbf{I} - \gamma_{M+i} \mathbf{H})^2 \mathbf{H} \right\rangle. \end{aligned}$$

Finally, applying Lemma C.2 completes the proof. $\square$

C.2 Lower Bounds

Lemma C.3 (Lower bounds on the variance iterates). *Suppose that Assumptions 1B and 2' hold. Suppose that $\gamma_0 < 1/\|\mathbf{G}\|_2$, $\gamma_M < 1/\|\mathbf{H}\|_2$. Let $M_{\text{eff}} := M/\log(M)$, $N_{\text{eff}} := N/\log(N)$.*

• For $\mathbb{J} := \{j \in \mathbb{N}_+ : \mu_j \geq 1/(M_{\text{eff}} \gamma_0)\}$, it holds that

$$\mathbf{C}_M \succeq \frac{\sigma^2}{400} \cdot \left(\frac{1}{M_{\text{eff}}} \cdot \mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}} \gamma_0^2 \cdot \mathbf{G}_{\mathbb{J}^c} \right).$$

• For $\mathbb{K} := \{k \in \mathbb{N}_+ : \lambda_k \geq 1/(N_{\text{eff}}\gamma_M)\}$, it holds that

$$\sum_{t=0}^{N-1} \gamma_{M+t}^2 \prod_{i=t+1}^{N-1} (\mathbf{I} - \gamma_{M+i} \mathbf{H})^2 \mathbf{H} \succeq \frac{1}{400} \cdot \left(\frac{1}{N_{\text{eff}}} \mathbf{H}_{\mathbb{K}}^{-1} + N_{\text{eff}} \gamma_M^2 \cdot \mathbf{H}_{\mathbb{K}^c} \right).$$

Proof. There are from the proof of Theorem 7 in Wu et al. [2021]. $\square$

Theorem C.2 (Variance error lower bound). *Suppose that Assumptions 1B and 2' hold. Suppose that $\gamma_0 < 1/\|\mathbf{G}\|_2$, $\gamma_M < 1/\|\mathbf{H}\|_2$. Let $M_{\text{eff}} := M/\log(M)$, $N_{\text{eff}} := N/\log(N)$. Then it holds that*

$$\langle \mathbf{H}, \mathbf{C}_{M+N} \rangle \geq \frac{\sigma^2}{400} \cdot \left(\frac{D_{\text{eff}}^{\text{finetune}}}{M_{\text{eff}}} + \frac{D_{\text{eff}}}{N_{\text{eff}}} \right),$$

where

$$D_{\text{eff}} := \text{tr}(\mathbf{H}\mathbf{H}_{\mathbb{K}}^{-1}) + N_{\text{eff}}^2 \gamma_M^2 \cdot \text{tr}(\mathbf{H}\mathbf{H}_{\mathbb{K}^c}),$$

$$D_{\text{eff}}^{\text{finetune}} := \text{tr} \left(\prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H})^2 \mathbf{H} \cdot (\mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}}^2 \gamma_0^2 \cdot \mathbf{G}_{\mathbb{J}^c}) \right),$$

and

$$\mathbb{K} := \{k \in \mathbb{N}_+ : \lambda_k \geq 1/(N_{\text{eff}}\gamma_M)\}, \quad \mathbb{J} := \{j \in \mathbb{N}_+ : \mu_j \geq 1/(M_{\text{eff}}\gamma_0)\}.$$

Proof of Theorem C.2. The proof idea is similar to that of Theorem C.1. For every $t = 0, \dots, N-1$, it holds that

$$\begin{aligned} \mathbf{C}_{M+t+1} &= (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{C}_{M+t} + \gamma_{M+t}^2 \sigma^2 \cdot \mathbf{H} \\ &\succeq (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{C}_{M+t} + \gamma_{M+t}^2 \sigma^2 \cdot \mathbf{H}. \end{aligned}$$

Solving the above recursion from $t = 0$ to $t = N-1$ we obtain

$$\begin{aligned} \mathbf{C}_{M+N} &\succeq \prod_{t=0}^{N-1} (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{C}_M + \sigma^2 \cdot \sum_{t=0}^{N-1} \gamma_{M+t}^2 \prod_{i=t+1}^{N-1} (\mathcal{I} - \gamma_{M+i} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+i})) \circ \mathbf{H} \\ &= \prod_{t=0}^{N-1} (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{C}_M + \sigma^2 \cdot \sum_{t=0}^{N-1} \gamma_{M+t}^2 \prod_{i=t+1}^{N-1} (\mathbf{I} - \gamma_{M+i} \mathbf{H})^2 \mathbf{H}. \end{aligned}$$

Therefore the variance error is

$$\begin{aligned} &\langle \mathbf{H}, \mathbf{C}_{M+N} \rangle \\ &\geq \left\langle \mathbf{H}, \prod_{t=0}^{N-1} (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{C}_M \right\rangle + \sigma^2 \left\langle \mathbf{H}, \sum_{t=0}^{N-1} \gamma_{M+t}^2 \prod_{i=t+1}^{N-1} (\mathbf{I} - \gamma_{M+i} \mathbf{H})^2 \mathbf{H} \right\rangle \\ &= \left\langle \prod_{t=0}^{N-1} (\mathcal{I} - \gamma_{M+t} \tilde{\mathcal{T}}_{\mathbf{H}}(\gamma_{M+t})) \circ \mathbf{H}, \mathbf{C}_M \right\rangle + \sigma^2 \left\langle \mathbf{H}, \sum_{t=0}^{N-1} \gamma_{M+t}^2 \prod_{i=t+1}^{N-1} (\mathbf{I} - \gamma_{M+i} \mathbf{H})^2 \mathbf{H} \right\rangle \\ &= \left\langle \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H})^2 \mathbf{H}, \mathbf{C}_M \right\rangle + \sigma^2 \left\langle \mathbf{H}, \sum_{t=0}^{N-1} \gamma_{M+t}^2 \prod_{i=t+1}^{N-1} (\mathbf{I} - \gamma_{M+i} \mathbf{H})^2 \mathbf{H} \right\rangle. \end{aligned}$$

Finally, applying Lemma C.3 completes the proof. $\square$

D Bias Error Analysis

D.1 Upper Bounds

Lemma D.1 (Bounds on the summation of bias iterates). *Suppose that Assumption 1A holds. Suppose that $\gamma < 1/(\alpha \text{tr}(\mathbf{G}))$. Then for every $n \geq 1$, it holds that*

$$\frac{1}{2\gamma} \cdot (\mathbf{I} - (\mathbf{I} - \gamma \mathbf{G})^{2n}) \preceq \sum_{t=0}^{n-1} (\mathcal{I} - \gamma \cdot \mathcal{T}_{\mathbf{G}}(\gamma))^t \circ \mathbf{G} \preceq \frac{1}{\gamma} \cdot (\mathbf{I} - (\mathbf{I} - \gamma \mathbf{G})^{2n}).$$

Proof. By definition and Assumption 1A, we have

$$\mathcal{T}_{\mathbf{G}}(\gamma) \circ \mathbf{I} = 2\mathbf{G} - \gamma \cdot \mathcal{M}_{\mathbf{G}} \circ \mathbf{I} \begin{cases} \preceq 2\mathbf{G}; \\ \succeq \mathbf{G}. \end{cases}$$

Notice that $\mathcal{T}_{\mathbf{G}}^{-1}(\gamma)$ is a PSD mapping (when operates on PSD matrices), therefore

$$\frac{1}{2} \cdot \mathbf{I} \preceq \mathcal{T}_{\mathbf{G}}^{-1}(\gamma) \circ \mathbf{G} \preceq \mathbf{I}.$$

Similarly, $\tilde{\mathcal{T}}_{\mathbf{G}}^{-1}(\gamma)$ is also a PSD mapping and that we have

$$\tilde{\mathcal{T}}_{\mathbf{G}}(\gamma) \circ \mathbf{I} = 2\mathbf{G} - \gamma \cdot \mathbf{G}^2 \begin{cases} \preceq 2\mathbf{G}; \\ \succeq \mathbf{G}, \end{cases}$$

therefore

$$\frac{1}{2} \cdot \mathbf{I} \preceq \tilde{\mathcal{T}}_{\mathbf{G}}^{-1}(\gamma) \circ \mathbf{G} \preceq \mathbf{I}.$$

With the above, we prove the upper bound as follows:

$$\begin{aligned} \sum_{t=0}^{n-1} (\mathcal{I} - \gamma \cdot \mathcal{T}_{\mathbf{G}}(\gamma))^t \circ \mathbf{G} &= \frac{1}{\gamma} \cdot \left(\mathcal{I} - (\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma))^n \right) \circ \mathcal{T}_{\mathbf{G}}^{-1}(\gamma) \circ \mathbf{G} \\ &\preceq \frac{1}{\gamma} \cdot \left(\mathcal{I} - (\mathcal{I} - \gamma \tilde{\mathcal{T}}_{\mathbf{G}}(\gamma))^n \right) \circ \mathcal{T}_{\mathbf{G}}^{-1}(\gamma) \circ \mathbf{G} \quad (\text{since } \tilde{\mathcal{T}} - \mathcal{T} \text{ is PSD}) \\ &\preceq \frac{1}{\gamma} \cdot \left(\mathcal{I} - (\mathcal{I} - \gamma \tilde{\mathcal{T}}_{\mathbf{G}}(\gamma))^n \right) \circ \mathbf{I} \quad (\text{since } \mathcal{T}_{\mathbf{G}}^{-1}(\gamma) \circ \mathbf{G} \preceq \mathbf{I}) \\ &= \frac{1}{\gamma} \cdot (\mathbf{I} - (\mathbf{I} - \gamma \mathbf{G})^{2n}). \end{aligned}$$

The lower bound is because

$$\begin{aligned} \sum_{t=0}^{n-1} (\mathcal{I} - \gamma \cdot \mathcal{T}_{\mathbf{G}}(\gamma))^t \circ \mathbf{G} &\succeq \sum_{t=0}^{n-1} (\mathcal{I} - \gamma \cdot \tilde{\mathcal{T}}_{\mathbf{G}}(\gamma))^t \circ \mathbf{G} \quad (\text{since } \tilde{\mathcal{T}} - \mathcal{T} \text{ is PSD}) \\ &= \frac{1}{\gamma} \cdot \left(\mathcal{I} - (\mathcal{I} - \gamma \tilde{\mathcal{T}}_{\mathbf{G}}(\gamma))^T \right) \circ \tilde{\mathcal{T}}_{\mathbf{G}}^{-1}(\gamma) \circ \mathbf{G} \\ &\succeq \frac{1}{2\gamma} \cdot \left(\mathcal{I} - (\mathcal{I} - \gamma \tilde{\mathcal{T}}_{\mathbf{G}}(\gamma))^T \right) \circ \mathbf{I} \quad (\text{since } \tilde{\mathcal{T}}_{\mathbf{G}}^{-1}(\gamma) \circ \mathbf{G} \succeq 0.5\mathbf{I}) \\ &= \frac{1}{2\gamma} \cdot (\mathbf{I} - (\mathbf{I} - \gamma \mathbf{G})^{2n}). \end{aligned}$$

□

Lemma D.2 (Crude bounds on the bias iterates). *Suppose that Assumption 1A holds. Suppose that $\gamma < 1/(2\alpha \text{tr}(\mathbf{G}))$. Then the following holds for every $n \geq 0$:*

$$(\mathcal{I} - \gamma \cdot \mathcal{T}_{\mathbf{G}}(\gamma))^n \circ \mathbf{G} \preceq \begin{cases} (1 + \alpha\gamma \text{tr}(\mathbf{G})) \cdot \mathbf{G}, \\ \frac{1}{1 - 2\alpha\gamma \text{tr}(\mathbf{G})} \cdot \frac{1}{\max\{n, 1\}\gamma} \cdot \mathbf{I}. \end{cases}$$

In particular, the following holds for every $n \geq 1$ and every index set $\mathbb{J} \subset \mathbb{N}_+$:

$$(\mathcal{I} - \gamma \cdot \mathcal{T}_{\mathbf{G}}(\gamma))^n \circ \mathbf{G} \preceq \frac{1}{1 - 2\alpha\gamma \text{tr}(\mathbf{G})} \cdot \left(\frac{\mathbf{I}_{\mathbb{J}}}{n\gamma} + \mathbf{G}_{\mathbb{J}^c} \right).$$

Proof. Notice the following decomposition:

$$\begin{aligned} &(\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma))^n \circ \mathbf{G} \\ &= (\mathcal{I} - \gamma \tilde{\mathcal{T}}_{\mathbf{G}}(\gamma))^n \circ \mathbf{G} + \gamma^2 \sum_{t=0}^{n-1} (\mathcal{I} - \gamma \tilde{\mathcal{T}}_{\mathbf{G}}(\gamma))^{n-1-t} \circ (\mathcal{M}_{\mathbf{G}} - \tilde{\mathcal{M}}_{\mathbf{G}}) \circ (\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma))^t \circ \mathbf{G} \end{aligned}$$

$$\begin{aligned}
&\preceq (\mathbf{I} - \gamma \mathbf{G})^{2n} \mathbf{G} + \alpha \gamma^2 \sum_{t=0}^{n-1} (\mathcal{I} - \gamma \tilde{\mathcal{T}}_{\mathbf{G}}(\gamma))^{n-1-t} \circ \mathbf{G} \cdot \langle \mathbf{G}, (\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma))^t \circ \mathbf{G} \rangle \\
&= (\mathbf{I} - \gamma \mathbf{G})^{2n} \mathbf{G} + \alpha \gamma^2 \sum_{t=0}^{n-1} (\mathbf{I} - \gamma \mathbf{G})^{2(n-1-t)} \mathbf{G} \cdot \langle \mathbf{G}, (\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma))^t \circ \mathbf{G} \rangle. \tag{10}
\end{aligned}$$

Based on (10) we show the first conclusion as follows:

$$\begin{aligned}
(\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma))^n \circ \mathbf{G} &\preceq \mathbf{G} + \alpha \gamma^2 \sum_{t=0}^{n-1} \mathbf{G} \cdot \langle \mathbf{G}, (\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma))^t \circ \mathbf{G} \rangle \\
&\preceq \mathbf{G} + \alpha \gamma^2 \mathbf{G} \cdot \langle \mathbf{G}, \frac{1}{\gamma} (\mathbf{I} - (\mathbf{I} - \gamma \mathbf{G})^{2n}) \rangle \quad (\text{by Lemma D.1}) \\
&\preceq \mathbf{G} + \alpha \gamma^2 \langle \mathbf{G}, \frac{1}{\gamma} \mathbf{I} \rangle = (1 + \alpha \gamma \text{tr}(\mathbf{G})) \cdot \mathbf{G}.
\end{aligned}$$

We now prove the second conclusion by induction. For $n = 1$, it holds because of Lemma D.1:

$$(\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma)) \circ \mathbf{G} \preceq \sum_{t=0}^1 (\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma))^t \circ \mathbf{G} \preceq \frac{1}{\gamma} \cdot \mathbf{I}.$$

Now consider $n \geq 2$ based on (10). We bound the second term in (10) separately for $\sum_{t=0}^{n/2-1}$ and $\sum_{t=n/2}^{n-1}$. For the first part,

$$\begin{aligned}
&\sum_{t=0}^{n/2-1} (\mathbf{I} - \gamma \mathbf{G})^{2(n-1-t)} \mathbf{G} \cdot \langle \mathbf{G}, (\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma))^t \circ \mathbf{G} \rangle \\
&\preceq (\mathbf{I} - \gamma \mathbf{G})^n \mathbf{G} \cdot \langle \mathbf{G}, \sum_{t=0}^{n/2-1} (\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma))^t \circ \mathbf{G} \rangle \\
&\preceq (\mathbf{I} - \gamma \mathbf{G})^n \mathbf{G} \cdot \langle \mathbf{G}, \frac{1}{\gamma} \mathbf{I} \rangle \quad (\text{by Lemma D.1}) \\
&\preceq \frac{\text{tr}(\mathbf{G})}{\gamma} \cdot (\mathbf{I} - \gamma \mathbf{G})^n \mathbf{G} \preceq \frac{\text{tr}(\mathbf{G})}{\gamma} \cdot \frac{1}{n\gamma} \cdot \mathbf{I}. \tag{11}
\end{aligned}$$

For the second part,

$$\begin{aligned}
&\sum_{t=n/2}^{n-1} (\mathbf{I} - \gamma \mathbf{G})^{2(n-1-t)} \mathbf{G} \cdot \langle \mathbf{G}, (\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma))^t \circ \mathbf{G} \rangle \\
&\preceq \sum_{t=n/2}^{n-1} (\mathbf{I} - \gamma \mathbf{G})^{2(n-1-t)} \mathbf{G} \cdot \langle \mathbf{G}, \frac{1}{1 - 2\alpha\gamma \text{tr}(\mathbf{G})} \cdot \frac{2}{n\gamma} \cdot \mathbf{I} \rangle \quad (\text{by induction hypothesis}) \\
&\preceq \frac{\text{tr}(\mathbf{G})}{1 - 2\alpha\gamma \text{tr}(\mathbf{G})} \cdot \frac{2}{n\gamma} \cdot \sum_{t=n/2}^{n-1} (\mathbf{I} - \gamma \mathbf{G})^{(n-1-t)} \mathbf{G} \\
&= \frac{\text{tr}(\mathbf{G})}{1 - 2\alpha\gamma \text{tr}(\mathbf{G})} \cdot \frac{2}{n\gamma} \cdot \frac{1}{\gamma} \cdot (\mathbf{I} - (\mathbf{I} - \gamma \mathbf{G})^{n/2}) \\
&\preceq \frac{\text{tr}(\mathbf{G})}{1 - 2\alpha\gamma \text{tr}(\mathbf{G})} \cdot \frac{2}{n\gamma} \cdot \frac{1}{\gamma} \cdot \mathbf{I}. \tag{12}
\end{aligned}$$

Inserting (11) and (12) into (10), and apply that $(\mathbf{I} - \gamma \mathbf{G})^{2n} \mathbf{G} \preceq \frac{1}{2n\gamma} \cdot \mathbf{I}$, we obtain that

$$\begin{aligned}
(\mathcal{I} - \gamma \mathcal{T}_{\mathbf{G}}(\gamma))^n \circ \mathbf{G} &\preceq \frac{1}{2n\gamma} \cdot \mathbf{I} + \frac{\alpha \text{tr}(\mathbf{G})}{n} \cdot \mathbf{I} + \frac{\alpha \text{tr}(\mathbf{G})}{1 - 2\alpha\gamma \text{tr}(\mathbf{G})} \cdot \frac{2}{n} \cdot \mathbf{I} \\
&= \left(\frac{1}{2} + \alpha \gamma \text{tr}(\mathbf{G}) + \frac{2\alpha\gamma \text{tr}(\mathbf{G})}{1 - 2\alpha\gamma \text{tr}(\mathbf{G})} \right) \cdot \frac{1}{n\gamma} \cdot \mathbf{I}
\end{aligned}$$

$$\preceq \frac{1}{1 - 2\alpha\gamma \operatorname{tr}(\mathbf{G})} \cdot \frac{1}{n\gamma} \cdot \mathbf{I}.$$

We have completed the induction. $\square$

Lemma D.3 (Bounds on the summation of bias iterates). *Suppose that Assumption 1A holds. Suppose that $\gamma_0 < 1/(2\alpha \operatorname{tr}(\mathbf{G}))$. Then the following holds for every index set $\mathbb{J} \subset \mathbb{N}_+$:*

$$\sum_{t=1}^{M_{\text{eff}}} \langle \mathbf{G}, \mathbf{B}_{t-1} \rangle \leq \frac{1}{\gamma_0} \cdot \langle \mathbf{I}_{\mathbb{J}} + 2M_{\text{eff}}\gamma_0 \mathbf{G}_{\mathbb{J}^c}, \mathbf{B}_0 \rangle.$$

Proof. Notice that

$$\sum_{t=1}^{M_{\text{eff}}} \langle \mathbf{G}, \mathbf{B}_{t-1} \rangle = \sum_{t=1}^{M_{\text{eff}}} \langle \mathbf{G}, (\mathcal{I} - \gamma_0 \mathcal{T}_{\mathbf{G}}(\gamma_0))^{t-1} \circ \mathbf{B}_0 \rangle = \langle \sum_{t=1}^{M_{\text{eff}}} (\mathcal{I} - \gamma_0 \mathcal{T}_{\mathbf{G}}(\gamma_0))^{t-1} \circ \mathbf{G}, \mathbf{B}_0 \rangle.$$

Then we apply Lemma D.1 to obtain that

$$\sum_{t=1}^{M_{\text{eff}}} (\mathcal{I} - \gamma_0 \cdot \mathcal{T}_{\mathbf{G}}(\gamma_0))^{t-1} \circ \mathbf{G} \preceq \frac{1}{\gamma_0} \cdot (\mathbf{I} - (\mathbf{I} - \gamma_0 \mathbf{G})^{2M_{\text{eff}}}) \preceq \frac{1}{\gamma_0} \cdot (\mathbf{I}_{\mathbb{J}} + 2M_{\text{eff}}\gamma_0 \mathbf{G}_{\mathbb{J}^c}).$$

This completes the proof. $\square$

Lemma D.4 (Crude bounds on the bias iterates). *Suppose that Assumption 1A holds. Suppose that $\gamma_0 < 1/(2\alpha \operatorname{tr}(\mathbf{G}))$. Then the following holds for every index set $\mathbb{J} \subset \mathbb{N}_+$ and $t \geq M_{\text{eff}}$:*

$$\langle \mathbf{G}, \mathbf{B}_t \rangle \leq \frac{e}{1 - 2\alpha \operatorname{tr}(\mathbf{G})\gamma_0} \cdot \left\langle \frac{1}{M_{\text{eff}}\gamma_0} \cdot \mathbf{I}_{0:j} + \mathbf{G}_{j:\infty}, \mathbf{B}_0 \right\rangle.$$

Proof. Let $L(t) = \lfloor t \log(M)/M \rfloor = \lfloor t/M_{\text{eff}} \rfloor$, then $L(t) \geq 1$ as $t \geq M_{\text{eff}}$. Notice that

$$\begin{aligned} \langle \mathbf{G}, \mathbf{B}_t \rangle &:= \left\langle \mathbf{G}, \prod_{t=0}^{t-1} (\mathcal{I} - \gamma_t \cdot \mathcal{T}_{\mathbf{G}}(\gamma_t)) \circ \mathbf{B}_0 \right\rangle \\ &= \left\langle \mathbf{G}, \left(\mathcal{I} - \frac{\gamma_0}{2^{L(t)}} \cdot \mathcal{T}_{\mathbf{G}}\left(\frac{\gamma_0}{2^{L(t)}}\right) \right)^{t-L(t)\log(M)} \circ \prod_{\ell=0}^{L(t)-1} \left(\mathcal{I} - \frac{\gamma_0}{2^\ell} \cdot \mathcal{T}_{\mathbf{G}}\left(\frac{\gamma_0}{2^\ell}\right) \right)^{M_{\text{eff}}} \circ \mathbf{B}_0 \right\rangle \\ &= \left\langle \prod_{\ell=L(t)-1}^0 \left(\mathcal{I} - \frac{\gamma_0}{2^\ell} \cdot \mathcal{T}_{\mathbf{G}}\left(\frac{\gamma_0}{2^\ell}\right) \right)^{M_{\text{eff}}} \circ \left(\mathcal{I} - \frac{\gamma_0}{2^{L(t)}} \cdot \mathcal{T}_{\mathbf{G}}\left(\frac{\gamma_0}{2^{L(t)}}\right) \right)^{t-L(t)\log(M)} \circ \mathbf{G}, \mathbf{B}_0 \right\rangle. \end{aligned}$$

We then recursively use Lemma D.2 to obtain that

$$\begin{aligned} &\prod_{\ell=L(t)-1}^0 \left(\mathcal{I} - \frac{\gamma_0}{2^\ell} \cdot \mathcal{T}_{\mathbf{G}}\left(\frac{\gamma_0}{2^\ell}\right) \right)^{M_{\text{eff}}} \circ \left(\mathcal{I} - \frac{\gamma_0}{2^{L(t)}} \cdot \mathcal{T}_{\mathbf{G}}\left(\frac{\gamma_0}{2^{L(t)}}\right) \right)^{t-L(t)\log(M)} \circ \mathbf{G} \\ &\preceq \left(1 + \frac{\gamma_0}{2^{L(t)}} \cdot \alpha \operatorname{tr}(\mathbf{G}) \right) \cdot \prod_{\ell=L(t)-1}^0 \left(\mathcal{I} - \frac{\gamma_0}{2^\ell} \cdot \mathcal{T}_{\mathbf{G}}\left(\frac{\gamma_0}{2^\ell}\right) \right)^{M_{\text{eff}}} \circ \mathbf{G} \\ &\preceq \prod_{\ell=1}^{L(t)} \left(1 + \frac{\gamma_0}{2^\ell} \cdot \alpha \operatorname{tr}(\mathbf{G}) \right) \cdot (\mathcal{I} - \gamma_0 \cdot \mathcal{T}_{\mathbf{G}}(\gamma_0))^{M_{\text{eff}}} \circ \mathbf{G} \\ &\preceq e^{\alpha \operatorname{tr}(\mathbf{G}) \sum_{\ell=1}^{L(t)} \gamma_0/2^\ell} \cdot (\mathcal{I} - \gamma_0 \cdot \mathcal{T}_{\mathbf{G}}(\gamma_0))^{M_{\text{eff}}} \circ \mathbf{G} \\ &\preceq e^{\alpha \operatorname{tr}(\mathbf{G})\gamma_0} \cdot (\mathcal{I} - \gamma_0 \cdot \mathcal{T}_{\mathbf{G}}(\gamma_0))^{M_{\text{eff}}} \circ \mathbf{G} \\ &\preceq e \cdot (\mathcal{I} - \gamma_0 \cdot \mathcal{T}_{\mathbf{G}}(\gamma_0))^{M_{\text{eff}}} \circ \mathbf{G} \preceq \frac{e}{1 - 2\alpha \operatorname{tr}(\mathbf{G})\gamma_0} \cdot \left(\frac{1}{M_{\text{eff}}\gamma_0} \cdot \mathbf{I}_{\mathbb{J}} + \mathbf{G}_{\mathbb{J}^c} \right). \end{aligned}$$

Combining everything we complete the proof. $\square$

Lemma D.5 (Upper bounds for the bias iterates). *Suppose that Assumption 1A holds. Suppose that $\gamma_0 < 1/(2\alpha \text{tr}(\mathbf{G}))$. Let $M_{\text{eff}} := M/\log(M)$. Then the following holds for every index set $\mathbb{J} \subset \mathbb{N}_+$:*

$$\begin{aligned} \mathbf{B}_M &\preceq \prod_{t=1}^M (\mathbf{I} - \gamma_t \mathbf{G}) \cdot \mathbf{B}_0 \cdot \prod_{t=1}^M (\mathbf{I} - \gamma_t \mathbf{G}) \\ &\quad + \frac{12e\alpha}{1 - 2\alpha \text{tr}(\mathbf{G})\gamma_0} \cdot \left\langle \frac{\mathbf{I}_{\mathbb{J}}}{M_{\text{eff}}\gamma_0} + \mathbf{G}_{\mathbb{J}^c}, \mathbf{B}_0 \right\rangle \cdot \frac{\mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}}^2 \gamma_0^2 \cdot \mathbf{G}_{\mathbb{J}^c}}{M_{\text{eff}}}. \end{aligned}$$

Proof. We begin with the following inequality:

$$\begin{aligned} \mathbf{B}_{t+1} &= (\mathcal{I} - \gamma_t \tilde{\mathcal{T}}_{\mathbf{G}}(\gamma_t)) \circ \mathbf{B}_t + \gamma_t^2 \cdot (\mathcal{M}_{\mathbf{G}} - \tilde{\mathcal{M}}_{\mathbf{G}}) \circ \mathbf{B}_t \\ &\preceq (\mathcal{I} - \gamma_t \tilde{\mathcal{T}}_{\mathbf{G}}(\gamma_t)) \circ \mathbf{B}_t + \alpha \gamma_t^2 \cdot \mathbf{G} \cdot \langle \mathbf{G}, \mathbf{B}_t \rangle, \end{aligned}$$

which implies that

$$\begin{aligned} \mathbf{B}_M &\preceq \prod_{t=0}^{M-1} (\mathcal{I} - \gamma_t \tilde{\mathcal{T}}_{\mathbf{G}}(\gamma_t)) \circ \mathbf{B}_0 + \alpha \cdot \sum_{t=0}^{M-1} \gamma_t^2 \cdot \prod_{i=t+1}^{M-1} (\mathcal{I} - \gamma_i \tilde{\mathcal{T}}_{\mathbf{G}}(\gamma_i)) \circ \mathbf{G} \cdot \langle \mathbf{G}, \mathbf{B}_t \rangle \\ &= \prod_{t=0}^{M-1} (\mathcal{I} - \gamma_t \tilde{\mathcal{T}}_{\mathbf{G}}(\gamma_t)) \circ \mathbf{B}_0 + \alpha \cdot \sum_{t=0}^{M-1} \gamma_t^2 \cdot \prod_{i=t+1}^M (\mathbf{I} - \gamma_i \mathbf{G})^2 \mathbf{G} \cdot \langle \mathbf{G}, \mathbf{B}_t \rangle. \end{aligned} \quad (13)$$

We next bound the second term in (13) separately for $\sum_{t=0}^{M_{\text{eff}}-1} (\cdot)$ and $\sum_{t=M_{\text{eff}}}^{M-1} (\cdot)$. For the first part,

$$\begin{aligned} &\sum_{t=0}^{M_{\text{eff}}-1} \gamma_t^2 \cdot \prod_{i=t+1}^M (\mathbf{I} - \gamma_i \mathbf{G})^2 \mathbf{G} \cdot \langle \mathbf{G}, \mathbf{B}_t \rangle = \gamma_0^2 \cdot \sum_{t=0}^{M_{\text{eff}}-1} \prod_{i=t+1}^M (\mathbf{I} - \gamma_i \mathbf{G})^2 \mathbf{G} \cdot \langle \mathbf{G}, \mathbf{B}_t \rangle \\ &\leq \gamma_0^2 \cdot \sum_{t=0}^{M_{\text{eff}}-1} \prod_{i=M_{\text{eff}}}^{2M_{\text{eff}}-1} (\mathbf{I} - \gamma_i \mathbf{G})^2 \mathbf{G} \cdot \langle \mathbf{G}, \mathbf{B}_t \rangle = \gamma_0^2 \cdot \left(\mathbf{I} - \frac{\gamma_0}{2} \mathbf{G} \right)^{2M_{\text{eff}}} \mathbf{G} \cdot \sum_{t=0}^{M_{\text{eff}}-1} \langle \mathbf{G}, \mathbf{B}_t \rangle \\ &\leq \gamma_0 \cdot \left(\mathbf{I} - \frac{\gamma_0}{2} \mathbf{G} \right)^{2M_{\text{eff}}} \mathbf{G} \cdot \langle \mathbf{I}_{\mathbb{J}} + 2M_{\text{eff}}\gamma_0 \mathbf{G}_{\mathbb{J}^c}, \mathbf{B}_0 \rangle \quad (\text{by Lemma D.3}) \\ &\leq \frac{8}{M_{\text{eff}}^2 \gamma_0} \cdot (\mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}}^2 \gamma_0^2 \mathbf{G}_{\mathbb{J}^c}) \cdot \langle \mathbf{I}_{\mathbb{J}} + M_{\text{eff}}\gamma_0 \mathbf{G}_{\mathbb{J}^c}, \mathbf{B}_0 \rangle, \end{aligned} \quad (14)$$

where in the last inequality we use that

$$\left(\mathbf{I} - \frac{\gamma_0}{2} \mathbf{G} \right)^{2M_{\text{eff}}} \preceq \left(\frac{2}{M_{\text{eff}}\gamma_0} \mathbf{G}_{\mathbb{J}}^{-1} + \mathbf{I}_{\mathbb{J}^c} \right)^2 \preceq 4 \cdot \left(\frac{1}{M_{\text{eff}}^2 \gamma_0^2} \mathbf{G}_{\mathbb{J}}^{-2} + \mathbf{I}_{\mathbb{J}^c} \right).$$

For the second part, we apply Lemma D.4 to obtain that

$$\begin{aligned} &\sum_{t=M_{\text{eff}}}^{M-1} \gamma_t^2 \cdot \prod_{i=t+1}^M (\mathbf{I} - \gamma_i \mathbf{G})^2 \mathbf{G} \cdot \langle \mathbf{G}, \mathbf{B}_t \rangle \\ &\leq \frac{e}{1 - 2\alpha \text{tr}(\mathbf{G})\gamma_0} \cdot \left\langle \frac{1}{M_{\text{eff}}\gamma_0} \cdot \mathbf{I}_{\mathbb{J}} + \mathbf{G}_{\mathbb{J}^c}, \mathbf{B}_0 \right\rangle \cdot \sum_{t=M_{\text{eff}}}^{M-1} \gamma_t^2 \cdot \prod_{i=t+1}^M (\mathbf{I} - \gamma_i \mathbf{G})^2 \mathbf{G} \\ &\leq \frac{8e}{1 - 2\alpha \text{tr}(\mathbf{G})\gamma_0} \cdot \left\langle \frac{1}{M_{\text{eff}}\gamma_0} \cdot \mathbf{I}_{\mathbb{J}} + \mathbf{G}_{\mathbb{J}^c}, \mathbf{B}_0 \right\rangle \cdot \left(\frac{1}{M_{\text{eff}}} \mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}} \gamma_0^2 \cdot \mathbf{G}_{\mathbb{J}^c} \right), \end{aligned} \quad (15)$$

where in the last inequality we use Lemma C.2 (by setting $\mathbf{H}$ to $\mathbf{G}$).

Finally, inserting (14) and (15) into (13) completes the proof. $\square$

Theorem D.1 (Bias error upper bound). *Suppose that Assumption 1A holds. Suppose that $\gamma_0 < \min\{1/(4\alpha \text{tr}(\mathbf{G})), 1/(\alpha \text{tr}(\mathbf{H}))\}$, $\gamma_M < 1/(4\alpha \text{tr}(\mathbf{H}))$. Let $M_{\text{eff}} := M/\log(M)$, $N_{\text{eff}} := N/\log(N)$. Then it holds that*

$$\langle \mathbf{H}, \mathbf{B}_{M+N} \rangle \leq \left\| \prod_{t=M}^{M+N-1} (\mathbf{I} - \gamma_t \mathbf{H}) \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2$$

$$\begin{aligned}
& + 24e\alpha \cdot \left\| \mathbf{w}_0 - \mathbf{w}^* \right\|_{\frac{\mathbf{I}_{\mathbb{J}}}{M_{\text{eff}}\gamma_0} + \mathbf{G}_{\mathbb{J}^c}}^2 \cdot \frac{D_{\text{eff}}^{\text{finetune}}}{M_{\text{eff}}} \\
& + 576e^2\alpha \cdot \left(\left\| \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\frac{\mathbf{I}_{\mathbb{K}}}{N_{\text{eff}}\gamma_M} + \mathbf{H}_{\mathbb{K}^c}}^2 + \left\| \mathbf{w}_0 - \mathbf{w}^* \right\|_{\frac{\mathbf{I}_{\mathbb{J}}}{M_{\text{eff}}\gamma_0} + \mathbf{G}_{\mathbb{J}^c}}^2 \right) \cdot \frac{D_{\text{eff}}}{N_{\text{eff}}},
\end{aligned}$$

where

$$\begin{aligned}
D_{\text{eff}} &:= \text{tr}(\mathbf{H}\mathbf{H}_{\mathbb{K}}^{-1}) + N_{\text{eff}}^2\gamma_M^2 \cdot \text{tr}(\mathbf{H}\mathbf{H}_{\mathbb{K}^c}), \\
D_{\text{eff}}^{\text{finetune}} &:= \text{tr} \left(\prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t}\mathbf{H})^2 \mathbf{H} \cdot (\mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}}^2\gamma_0^2 \cdot \mathbf{G}_{\mathbb{J}^c}) \right),
\end{aligned}$$

and $\mathbb{K}, \mathbb{J}$ can be arbitrary index sets.

Proof. We first apply Lemma D.5 by setting $\mathbf{B}_0$ to $\mathbf{B}_M$, γ_0 to γ_M , M to N and $\mathbf{G}$ to $\mathbf{H}$, so that we have

$$\begin{aligned}
\mathbf{B}_{M+N} &\preceq \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t}\mathbf{H})\mathbf{B}_M \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t}\mathbf{H}) \\
&+ 24e\alpha \cdot \left\langle \frac{\mathbf{I}_{\mathbb{K}}}{N_{\text{eff}}\gamma_M} + \mathbf{H}_{\mathbb{K}^c}, \mathbf{B}_M \right\rangle \cdot \frac{\mathbf{H}_{\mathbb{K}}^{-1} + N_{\text{eff}}^2\gamma_M^2 \cdot \mathbf{H}_{\mathbb{K}^c}}{N_{\text{eff}}}.
\end{aligned}$$

Taking inner product with $\mathbf{H}$ we obtain

$$\begin{aligned}
\langle \mathbf{H}, \mathbf{B}_{M+N} \rangle &\preceq \left\langle \mathbf{H}, \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t}\mathbf{H})\mathbf{B}_M \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t}\mathbf{H}) \right\rangle \\
&+ 24e\alpha \cdot \left\langle \frac{\mathbf{I}_{\mathbb{K}}}{N_{\text{eff}}\gamma_M} + \mathbf{H}_{\mathbb{K}^c}, \mathbf{B}_M \right\rangle \cdot \frac{D_{\text{eff}}}{N_{\text{eff}}} \\
&= \left\langle \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t}\mathbf{H})^2 \mathbf{H}, \mathbf{B}_M \right\rangle + 24e\alpha \cdot \left\langle \frac{\mathbf{I}_{\mathbb{K}}}{N_{\text{eff}}\gamma_M} + \mathbf{H}_{\mathbb{K}^c}, \mathbf{B}_M \right\rangle \cdot \frac{D_{\text{eff}}}{N_{\text{eff}}}.
\end{aligned}$$

Now applying the upper bound for $\mathbf{B}_M$ in Lemma D.5, we obtain

$$\begin{aligned}
&\langle \mathbf{H}, \mathbf{B}_{M+N} \rangle \\
&\preceq \underbrace{\left\langle \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t}\mathbf{H})^2 \mathbf{H}, \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})\mathbf{B}_0 \prod_{t=1}^M (\mathbf{I} - \gamma_t \mathbf{G}) \right\rangle}_{(\spadesuit)} \\
&+ 24e\alpha \cdot \underbrace{\left\langle \frac{\mathbf{I}_{\mathbb{J}}}{M_{\text{eff}}\gamma_0} + \mathbf{G}_{\mathbb{J}^c}, \mathbf{B}_0 \right\rangle}_{(\heartsuit)} \cdot \frac{D_{\text{eff}}^{\text{finetune}}}{M_{\text{eff}}} \\
&+ 24e\alpha \cdot \underbrace{\left\langle \frac{\mathbf{I}_{\mathbb{K}}}{N_{\text{eff}}\gamma_M} + \mathbf{H}_{\mathbb{K}^c}, \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})\mathbf{B}_0 \prod_{t=1}^M (\mathbf{I} - \gamma_t \mathbf{G}) \right\rangle}_{(\diamond)} \cdot \frac{D_{\text{eff}}}{N_{\text{eff}}} \\
&+ 576e^2\alpha \cdot \underbrace{\left\langle \frac{\mathbf{I}_{\mathbb{J}}}{M_{\text{eff}}\gamma_0} + \mathbf{G}_{\mathbb{J}^c}, \mathbf{B}_0 \right\rangle}_{(\heartsuit)} \cdot \underbrace{\alpha \left\langle \frac{\mathbf{I}_{\mathbb{K}}}{N_{\text{eff}}\gamma_M} + \mathbf{H}_{\mathbb{K}^c}, \frac{\mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}}^2\gamma_0^2 \cdot \mathbf{G}_{\mathbb{J}^c}}{M_{\text{eff}}} \right\rangle}_{(\clubsuit)} \cdot \frac{D_{\text{eff}}}{N_{\text{eff}}}.
\end{aligned}$$

By definition we know that

$$\begin{aligned}
(\spadesuit) &= \left\| \prod_{t=M}^{M+N-1} (\mathbf{I} - \gamma_t \mathbf{H}) \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2, \\
(\heartsuit) &= \left\| \mathbf{w}_0 - \mathbf{w}^* \right\|_{\frac{\mathbf{I}_{\mathbb{J}}}{M_{\text{eff}}\gamma_0} + \mathbf{G}_{\mathbb{J}^c}}^2,
\end{aligned}$$

$$(\diamond) = \left\| \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\frac{\mathbf{I}_{\mathbb{K}}}{N_{\text{eff}} \gamma_M} + \mathbf{H}_{\mathbb{K}^c}}^2.$$

As for ($\clubsuit$), we can choose $\mathbb{K} = \emptyset$ and $\mathbb{J} = \{j : \mu_j \geq 1/(M_{\text{eff}} \gamma_0)\}$ so that

$$(\clubsuit) \leq \alpha \langle \mathbf{H}, \gamma_0 \mathbf{I} \rangle \leq \alpha \gamma_0 \text{tr}(\mathbf{H}) \leq 1.$$

Putting everything together completes the proof. $\square$

D.2 Lower Bounds

Lemma D.6 (Lower bounds for the bias iterates). *Suppose that Assumption 1B holds. Suppose that $\gamma_0 < 1/\|\mathbf{G}\|_2$. Let $M_{\text{eff}} := M/\log(M)$ and suppose that $M_{\text{eff}} \geq 10$. Let $\mathbb{J} := \{j : \mu_j \geq 1/(M_{\text{eff}} \gamma_0)\}$. Then it holds that*

$$\mathbf{B}_M \succeq \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G}) \mathbf{B}_0 \prod_{t=1}^M (\mathbf{I} - \gamma_t \mathbf{G}) + \frac{\beta}{1200} \cdot \langle \mathbf{G}_{\mathbb{J}^c}, \mathbf{B}_0 \rangle \cdot \frac{\mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}}^2 \gamma_0^2 \cdot \mathbf{G}_{\mathbb{J}^c}}{M_{\text{eff}}}.$$

Proof. This is from Theorem 8 in Wu et al. [2021]. $\square$

Theorem D.2 (Lower bounds for the bias error). *Suppose that Assumption 1B holds. Suppose that $\gamma_0 < 1/\|\mathbf{G}\|_2$, $\gamma_M < 1/\|\mathbf{H}\|_2$. Let $M_{\text{eff}} := M/\log(M)$, $N_{\text{eff}} := N/\log(N)$, and suppose that $M_{\text{eff}}, N_{\text{eff}} \geq 10$. Let $\mathbb{J} := \{j : \mu_j \geq 1/(M_{\text{eff}} \gamma_0)\}$, $\mathbb{K} := \{k : \lambda_k \geq 1/(N_{\text{eff}} \gamma_M)\}$. Then it holds that*

$$\begin{aligned} \langle \mathbf{H}, \mathbf{B}_{M+N} \rangle &\geq \left\| \prod_{t=M}^{M+N-1} (\mathbf{I} - \gamma_t \mathbf{H}) \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 \\ &+ \frac{\beta}{1200} \cdot \|\mathbf{w}_0 - \mathbf{w}^*\|_{\mathbf{G}_{\mathbb{J}^c}}^2 \cdot \frac{D_{\text{eff}}^{\text{finetune}}}{M_{\text{eff}}} + \frac{\beta}{1200} \cdot \left\| \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}_{\mathbb{K}^c}}^2 \cdot \frac{D_{\text{eff}}}{N_{\text{eff}}}, \end{aligned}$$

where

$$\begin{aligned} D_{\text{eff}} &:= \text{tr}(\mathbf{H} \mathbf{H}_{\mathbb{K}}^{-1}) + N_{\text{eff}}^2 \gamma_M^2 \cdot \text{tr}(\mathbf{H} \mathbf{H}_{\mathbb{K}^c}), \\ D_{\text{eff}}^{\text{finetune}} &:= \text{tr} \left(\prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H})^2 \mathbf{H} \cdot (\mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}}^2 \gamma_0^2 \cdot \mathbf{G}_{\mathbb{J}^c}) \right). \end{aligned}$$

Proof. We first apply Lemma D.6 by setting $\mathbf{B}_0$ to $\mathbf{B}_M$, γ_0 to γ_M , M to N and $\mathbf{G}$ to $\mathbf{H}$, so that we have

$$\mathbf{B}_{M+N} \succeq \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H}) \mathbf{B}_M \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H}) + \frac{\beta}{1200} \langle \mathbf{H}_{\mathbb{K}^c}, \mathbf{B}_M \rangle \frac{\mathbf{H}_{\mathbb{K}}^{-1} + N_{\text{eff}}^2 \gamma_M^2 \cdot \mathbf{H}_{\mathbb{K}^c}}{N_{\text{eff}}}.$$

Taking inner product with $\mathbf{H}$ we obtain

$$\begin{aligned} \langle \mathbf{H}, \mathbf{B}_{M+N} \rangle &\succeq \left\langle \mathbf{H}, \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H}) \mathbf{B}_M \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H}) \right\rangle + \frac{\beta}{1200} \cdot \langle \mathbf{H}_{\mathbb{K}^c}, \mathbf{B}_M \rangle \cdot \frac{D_{\text{eff}}}{N_{\text{eff}}} \\ &= \left\langle \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H})^2 \mathbf{H}, \mathbf{B}_M \right\rangle + \frac{\beta}{1200} \cdot \langle \mathbf{H}_{\mathbb{K}^c}, \mathbf{B}_M \rangle \cdot \frac{D_{\text{eff}}}{N_{\text{eff}}}. \end{aligned}$$

Now applying the lower bound for $\mathbf{B}_M$ in Lemma D.6, we obtain

$$\begin{aligned} \langle \mathbf{H}, \mathbf{B}_{M+N} \rangle &\succeq \left\langle \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H})^2 \mathbf{H}, \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G}) \mathbf{B}_0 \prod_{t=1}^M (\mathbf{I} - \gamma_t \mathbf{G}) \right\rangle \\ &+ \frac{\beta}{1200} \cdot \langle \mathbf{G}_{\mathbb{J}^c}, \mathbf{B}_0 \rangle \cdot \left\langle \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H})^2 \mathbf{H}, \frac{\mathbf{G}_{\mathbb{J}}^{-1} + M_{\text{eff}}^2 \gamma_0^2 \cdot \mathbf{G}_{\mathbb{J}^c}}{M_{\text{eff}}} \right\rangle \end{aligned}$$

$$\begin{aligned}
& + \frac{\beta}{1200} \cdot \left\langle \mathbf{H}_{\mathbb{K}^c}, \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G}) \mathbf{B}_0 \prod_{t=1}^M (\mathbf{I} - \gamma_t \mathbf{G}) \right\rangle \cdot \frac{D_{\text{eff}}}{N_{\text{eff}}} \\
& = \left\| \prod_{t=M}^{M+N-1} (\mathbf{I} - \gamma_t \mathbf{H}) \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G}) (\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 \\
& + \frac{\beta}{1200} \cdot \|\mathbf{w}_0 - \mathbf{w}^*\|_{\mathbf{G}_{\mathbb{J}^c}}^2 \cdot \frac{D_{\text{eff}}^{\text{finetune}}}{M_{\text{eff}}} \\
& + \frac{\beta}{1200} \cdot \left\| \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G}) (\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}_{\mathbb{K}^c}}^2 \cdot \frac{D_{\text{eff}}}{N_{\text{eff}}},
\end{aligned}$$

which completes the proof. $\square$

E Proof of Theorems in Main Text

E.1 Proof of Theorem 3.1

Proof of Theorem 3.1. This is by combining Theorems C.1 and D.1. $\square$

E.2 Proof of Theorem 3.2

Proof of Theorem 3.2. This is by combining Theorems C.2 and D.2. $\square$

E.3 Proof of Theorem 4.1

Proof of Theorem 4.1. During the proof, we use γ^{sl} and γ_0 to denote the initial stepsizes for supervised learning and pretraining, respectively. Then Corollaries 3.4 and 3.3 sharply characterize the risk bounds for supervised learning and pretraining, respectively. In particular, let $\text{SNR} := \alpha \|\mathbf{w}^*\|_{\mathbf{G}}^2 / \sigma^2$, then we have

$$\text{ExcessRisk}(\mathbf{w}_{0+N^{\text{sl}}}) \gtrsim \left\| \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_t^{\text{sl}} \mathbf{H}) (\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 + \sigma^2 \cdot \frac{D_{\text{eff}}^{\text{sl}}}{N_{\text{eff}}^{\text{sl}}}, \quad (16)$$

$$\text{ExcessRisk}(\mathbf{w}_{M+0}) \lesssim \left\| \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G}) (\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 + (1 + \text{SNR}) \sigma^2 \cdot \frac{D_{\text{eff}}^{\text{pretrain}}}{M_{\text{eff}}}. \quad (17)$$

Fix hyperparameters $(N_{\text{eff}}^{\text{sl}}, \gamma^{\text{sl}})$ for supervised learning, we now identify hyperparameters $(M_{\text{eff}}, \gamma_0)$ for pretraining so that its risk (17) is no larger than that of supervised learning (16) upto a constant factor. To this end, we claim that

$$\frac{D_{\text{eff}}^{\text{pretrain}}}{M_{\text{eff}}} \leq \frac{D_{\text{eff}}^{\text{sl}}}{N_{\text{eff}}^{\text{sl}}} \text{ given that } \gamma_0 \leq \frac{D_{\text{eff}}^{\text{sl}}}{N_{\text{eff}}^{\text{sl}} \text{tr}(\mathbf{H})}. \quad (18)$$

$$\begin{aligned}
\left\| \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G}) (\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 & \lesssim \left\| \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_t^{\text{sl}} \mathbf{H}) (\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 \\
& \text{given that } M_{\text{eff}} \gamma_0 \geq 4 N_{\text{eff}}^{\text{sl}} \gamma^{\text{sl}} \|\mathbf{H}_{0:k^*}\|_{\mathbf{G}}.
\end{aligned} \quad (19)$$

To prove (18), we consider the optimal index set $\mathbb{J}^* := \{j : \mu_j \geq 1/(M_{\text{eff}} \gamma_0)\}$ as defined in Corollary 3.3, then by definition we have

$$\begin{aligned}
\frac{D_{\text{eff}}^{\text{pretrain}}}{M_{\text{eff}}} & \leq \frac{1}{M_{\text{eff}}} \cdot \sum_{j \in \mathbb{J}^*} \mathbf{H}_{jj} \cdot \frac{1}{\mu_j} + M_{\text{eff}} \gamma_0 \cdot \sum_{j \notin \mathbb{J}^*} \mathbf{H}_{jj} \cdot \mu_j \\
& \leq \gamma_0 \cdot \sum_{j \in \mathbb{J}^*} \mathbf{H}_{jj} + \gamma_0 \cdot \sum_{j \notin \mathbb{J}^*} \mathbf{H}_{jj} = \gamma_0 \text{tr}(\mathbf{H}),
\end{aligned}$$

which justifies (18).

To prove (19), we consider the bias error separately in its head part and its tail part, divided by the optimal index $k^* := \min\{k : \lambda_k \geq 1/(N_{\text{eff}}^{\text{sl}} \gamma^{\text{sl}})\}$ as defined in Corollary 3.4. For a tail index $k > k^*$, we have

$$\prod_{t=0}^{M-1} (1 - \gamma_t \mu_k)^2 \leq 1 \leq 100 \cdot (1 - 2\gamma^{\text{sl}} \lambda_k)^{2N_{\text{eff}}^{\text{sl}}} \leq 100 \cdot \prod_{t=0}^{N_{\text{eff}}^{\text{sl}}-1} (1 - \gamma_t^{\text{sl}} \lambda_k)^2, \quad (20)$$

where the second inequality is because that $\gamma^{\text{sl}} \lambda_k \geq 1/N_{\text{eff}}^{\text{sl}}$ and $N_{\text{eff}}^{\text{sl}} \geq 10$. For a head index $k \leq k^*$, we have

$$\prod_{t=0}^{M-1} (1 - \gamma_t \mu_k)^2 \leq (1 - \gamma_0 \mu_k)^{2M_{\text{eff}}} \leq (1 - 2\gamma^{\text{sl}} \lambda_k)^{2N_{\text{eff}}^{\text{sl}}} \leq \prod_{t=0}^{N_{\text{eff}}^{\text{sl}}-1} (1 - \gamma_t^{\text{sl}} \lambda_k)^2, \quad (21)$$

where the second inequality is because:

$$\begin{aligned} (1 - \gamma_0 \mu_k)^{\frac{M_{\text{eff}}}{N_{\text{eff}}^{\text{sl}}}} &\leq 1 - \frac{M_{\text{eff}}}{2N_{\text{eff}}^{\text{sl}}} \cdot \gamma_0 \mu_k \quad (\text{since } ab \leq (a+b)/2 \text{ for } 0 < a, b < 1) \\ &\leq 1 - \frac{M_{\text{eff}}}{2N_{\text{eff}}^{\text{sl}} \|\mathbf{H}_{0:k^*}\|_{\mathbf{G}}} \cdot \gamma_0 \lambda_k \quad (\text{since } \|\mathbf{H}_{0:k^*}\|_{\mathbf{G}} := \max\{\lambda_k / \mu_k : k \leq k^*\}) \\ &\leq 1 - 2\gamma^{\text{sl}} \lambda_k. \quad (\text{use the condition in (19)}) \end{aligned}$$

Combining (20), (21) and the definition of bias error justifies (19).

Finally, we choose $\gamma_0 = D_{\text{eff}}^{\text{sl}} / (N_{\text{eff}}^{\text{sl}} \text{tr}(\mathbf{H}))$ and

$$M_{\text{eff}} \geq (N_{\text{eff}}^{\text{sl}})^2 \cdot \frac{4\|\mathbf{H}_{0:k^*}\|_{\mathbf{G}}}{\alpha D_{\text{eff}}} \geq (N_{\text{eff}}^{\text{sl}})^2 \cdot \frac{4\gamma^{\text{sl}} \text{tr}(\mathbf{H}) \|\mathbf{H}_{0:k^*}\|_{\mathbf{G}}}{D_{\text{eff}}^{\text{sl}}} = \frac{4N_{\text{eff}}^{\text{sl}} \gamma^{\text{sl}} \|\mathbf{H}_{0:k^*}\|_{\mathbf{G}}}{\gamma_0},$$

so that both (18) and (21) hold, which imply that the risk of pretraining (17) is no larger than of supervised learning (16) upto a constant factor. $\square$

E.4 Proof of Theorem 4.2

Proof of Theorem 4.2. During the proof, we use γ^{sl} , γ_0 and γ_M to denote the initial stepsizes for supervised learning, pretraining and finetuning, respectively. Then Corollary 3.4 and Theorem 3.1 sharply characterize the risk bounds for supervised learning and pretraining-finetuning, respectively. In particular, let $\text{SNR} := \alpha(\|\mathbf{w}^*\|_{\mathbf{G}}^2 + \|\mathbf{w}^*\|_{\mathbf{H}}^2) / \sigma^2$, then we have the following upper bound for pretraining-finetuning:

$$\begin{aligned} \text{ExcessRisk}(\mathbf{w}_{M+N}) &\lesssim \left\| \prod_{t=M}^{M+N-1} (\mathbf{I} - \gamma_t \mathbf{H}) \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 \\ &\quad + (1 + \text{SNR})\sigma^2 \cdot \left(\frac{D_{\text{eff}}^{\text{finetune}}}{M_{\text{eff}}} + \frac{D_{\text{eff}}}{N_{\text{eff}}} \right), \end{aligned} \quad (22)$$

and we have a lower bound for supervised learning shown in (16). Fix hyperparameters $(N_{\text{eff}}^{\text{sl}}, \gamma^{\text{sl}})$ for supervised learning, we now identify hyperparameters $(M_{\text{eff}}, N_{\text{eff}}, \gamma_0, \gamma_M)$ for pretraining-finetuning so that its risk (22) is no larger than that of supervised learning (16) upto a constant factor. To this end, we claim that

$$\frac{D_{\text{eff}}}{N_{\text{eff}}} \leq \frac{D_{\text{eff}}^{\text{sl}}}{N_{\text{eff}}^{\text{sl}}} \text{ given that } \gamma_M \leq \frac{D_{\text{eff}}^{\text{sl}}}{N_{\text{eff}}^{\text{sl}} \text{tr}(\mathbf{H})} \quad (23)$$

$$\frac{D_{\text{eff}}^{\text{finetune}}}{M_{\text{eff}}} \leq \frac{D_{\text{eff}}^{\text{sl}}}{N_{\text{eff}}^{\text{sl}}} \text{ given that } \gamma_0 \leq \frac{D_{\text{eff}}^{\text{sl}}}{N_{\text{eff}}^{\text{sl}} \text{tr}(\prod_{t=0}^{N-1} (\mathbf{I} - \gamma_{M+t} \mathbf{H})^2 \mathbf{H})}. \quad (24)$$

$$\begin{aligned} \left\| \prod_{t=M}^{M+N-1} (\mathbf{I} - \gamma_t \mathbf{H}) \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G}) \mathbf{w}^* \right\|_{\mathbf{H}}^2 &\lesssim \left\| \prod_{t=0}^{N-1} (\mathbf{I} - \gamma_t^{\text{sl}} \mathbf{H}) \mathbf{w}^* \right\|_{\mathbf{H}}^2 + \|\mathbf{w}^*\|_{\mathbf{H}}^2 \cdot \frac{D_{\text{eff}}^{\text{sl}}}{N_{\text{eff}}^{\text{sl}}} \\ &\text{given that } M_{\text{eff}} \gamma_0 \geq 4N_{\text{eff}}^{\text{sl}} \gamma^{\text{sl}} \|\mathbf{H}_{k^*:k^*}\|_{\mathbf{G}}. \end{aligned} \quad (25)$$

To prove (23) and (24), one only needs to repeat the proof for (18).

To prove (25), we consider the bias error separately in its head part, middle part and tail part, divided by the index

$$k^\dagger := \{k : \lambda_k \geq \log(N_{\text{eff}}^{\text{s1}})/(N_{\text{eff}}\gamma_M)\}$$

and the index $k^* := \min\{k : \lambda_k \geq 1/(N_{\text{eff}}^{\text{s1}}\gamma^{\text{s1}})\}$ as defined in Corollary 3.4. For a tail index $k > k^*$, we have

$$\prod_{t=M}^{M+N-1} (1-\gamma_t\lambda_k)^2 \prod_{t=0}^{M-1} (1-\gamma_t\mu_k)^2 \leq 1 \leq 100 \cdot (1-2\gamma^{\text{s1}}\lambda_k)^{2N_{\text{eff}}^{\text{s1}}} \leq 100 \cdot \prod_{t=0}^{N^{\text{s1}}-1} (1-\gamma_t^{\text{s1}}\lambda_k)^2, \quad (26)$$

where the second inequality is because that $\gamma^{\text{s1}}\lambda_k \geq 1/N_{\text{eff}}^{\text{s1}}$ and $N_{\text{eff}}^{\text{s1}} \geq 10$. For a middle index $k^\dagger < k \leq k^*$, we have

$$\prod_{t=M}^{M+N-1} (1-\gamma_t\lambda_k)^2 \prod_{t=0}^{M-1} (1-\gamma_t\mu_k)^2 \leq (1-\gamma_0\mu_k)^{2M_{\text{eff}}} \leq (1-2\gamma^{\text{s1}}\lambda_k)^{2N_{\text{eff}}^{\text{s1}}} \leq \prod_{t=0}^{N^{\text{s1}}-1} (1-\gamma_t^{\text{s1}}\lambda_k)^2, \quad (27)$$

where the second inequality is because:

$$\begin{aligned} (1-\gamma_0\mu_k)^{\frac{M_{\text{eff}}}{N_{\text{eff}}^{\text{s1}}}} &\leq 1 - \frac{M_{\text{eff}}}{2N_{\text{eff}}^{\text{s1}}} \cdot \gamma_0\mu_k \quad (\text{since } ab \leq (a+b)/2 \text{ for } 0 < a, b < 1) \\ &\leq 1 - \frac{M_{\text{eff}}\gamma_0\lambda_k}{2N_{\text{eff}}^{\text{s1}}\|\mathbf{H}_{k^\dagger:k^*}\|_{\mathbf{G}}} \quad (\text{since } \|\mathbf{H}_{k^\dagger:k^*}\|_{\mathbf{G}} := \max\{\lambda_k/\mu_k : k^\dagger < k \leq k^*\}) \\ &\leq 1 - 2\gamma^{\text{s1}}\lambda_k. \quad (\text{use the condition in (25)}) \end{aligned}$$

For a head index $k \leq k^\dagger$, we have

$$\prod_{t=M}^{M+N-1} (1-\gamma_t\lambda_k)^2 \prod_{t=0}^{M-1} (1-\gamma_t\mu_k)^2 \leq (1-\gamma_M\lambda_k)^{2N_{\text{eff}}} \leq e^{-2N_{\text{eff}}\gamma_M\lambda_k} \leq \frac{D_{\text{eff}}^{\text{s1}}}{N_{\text{eff}}^{\text{s1}}}, \quad (28)$$

where in the last inequality we use $\lambda_k \geq \lambda_{k^\dagger} \geq \log(N_{\text{eff}}^{\text{s1}})/(N_{\text{eff}}\gamma_M)$.

Combining (26), (27) (28) and the definition of bias error justifies (25).

Finally, we choose

$$\gamma_0 = \gamma_M = D_{\text{eff}}^{\text{s1}}/(N_{\text{eff}}^{\text{s1}} \text{tr}(\mathbf{H})),$$

$$k^\dagger := \{k : \lambda_k \geq \log(N_{\text{eff}}^{\text{s1}})/(N_{\text{eff}}\gamma_M) = N_{\text{eff}}^{\text{s1}} \log(N_{\text{eff}}^{\text{s1}}) \text{tr}(\mathbf{H})/(N_{\text{eff}}D_{\text{eff}}^{\text{s1}})\},$$

and

$$M_{\text{eff}} \geq (N_{\text{eff}}^{\text{s1}})^2 \cdot \frac{4\|\mathbf{H}_{k^\dagger:k^*}\|_{\mathbf{G}}}{\alpha D_{\text{eff}}^{\text{s1}}} \geq (N_{\text{eff}}^{\text{s1}})^2 \cdot \frac{4\gamma^{\text{s1}} \text{tr}(\mathbf{H})\|\mathbf{H}_{k^\dagger:k^*}\|_{\mathbf{G}}}{D_{\text{eff}}^{\text{s1}}} = \frac{4N_{\text{eff}}^{\text{s1}}\gamma^{\text{s1}}\|\mathbf{H}_{k^\dagger:k^*}\|_{\mathbf{G}}}{\gamma_0},$$

so that all (23), (24) and (28) hold, which imply that the risk of pretraining (22) is no larger than of supervised learning (16) upto a constant factor. $\square$

E.5 Proof of Example 4.3

Proof of Example 4.3. One may verify that $\text{tr}(\mathbf{H}) \approx \text{tr}(\mathbf{G}) \approx 1$ and that $\|\mathbf{w}^*\|_{\mathbf{H}}^2 \approx \|\mathbf{w}^*\|_{\mathbf{G}}^2 \approx \sigma^2 \approx 1$. Therefore

$$\gamma_0 \lesssim 1/(\text{tr}(\mathbf{H})) \approx 1, \gamma_M \lesssim 1/(\text{tr}(\mathbf{G})) \approx 1.$$

Pretraining. From Corollary 3.3, we know that

$$\text{ExcessRisk}(\mathbf{w}_{M+0}) \geq \lambda_1(1-2\gamma_0\mu_1)^{2M_{\text{eff}}}(\mathbf{w}^*[1])^2 \geq (1-2\gamma_0\epsilon^2)^{2M_{\text{eff}}},$$

for which to be smaller than ϵ one has to set

$$M_{\text{eff}} \gtrsim \gamma_0^{-1}\epsilon^{-2} \log \epsilon^{-1} \gtrsim \epsilon^{-2}.$$

Supervised Learning. As for supervised learning, we discuss its rate based on Corollary 3.4 and the choice of $\mathbb{K} = \{k : \lambda_k \geq 1/(N_{\text{eff}}\gamma_0)\}$.

- If $|\mathbb{K}| > \epsilon^{-0.5}$, then by Corollary 3.4 we have

$$\text{ExcessRisk}(\mathbf{w}_{0+N}) \gtrsim \sigma^2 \cdot \frac{|\mathbb{K}|}{N_{\text{eff}}} \gtrsim \frac{\epsilon^{-0.5}}{N_{\text{eff}}},$$

for which to be smaller than ϵ one has to have $N_{\text{eff}} \gtrsim \epsilon^{-1.5}$.

- If $|\mathbb{K}| \leq \epsilon^{-0.5}$, then by Corollary 3.4 we have

$$\text{ExcessRisk}(\mathbf{w}_{0+N}) \gtrsim \sigma^2 \cdot N_{\text{eff}}\gamma_0^2 \sum_{i>k^*} \lambda_i^2 \gtrsim N_{\text{eff}}\gamma_0^2 \epsilon^{0.5},$$

for which to be smaller than ϵ one has to have

$$N_{\text{eff}}\gamma_0^2 \lesssim \epsilon^{0.5}. \quad (29)$$

On the other hand, by Corollary 3.4 we have

$$\text{ExcessRisk}(\mathbf{w}_{0+N}) \geq \lambda_2(1 - 2\gamma_0\lambda_2)^{2N_{\text{eff}}}(\mathbf{w}^*[2])^2 \geq \epsilon^{0.5} \cdot (1 - 2\gamma_0\epsilon^{0.5})^{2N_{\text{eff}}},$$

for which to be smaller than ϵ one need to set

$$N_{\text{eff}}\gamma_0 \gtrsim \epsilon^{-0.5} \log \epsilon^{-0.5} \gtrsim \epsilon^{-0.5}. \quad (30)$$

Then (29) and (30) together imply that $N_{\text{eff}} \gtrsim \epsilon^{-1.5}$.

In sum, for $\text{ExcessRisk}(\mathbf{w}_{0+N}) \leq \epsilon$ one has to set

$$N_{\text{eff}} \gtrsim \epsilon^{-1.5}.$$

Pretraining-Finetuning. Now we consider pretraining-finetuning by Theorem 3.1. We set

$$\gamma_0 \approx 1, \quad \gamma_M \approx \epsilon, \quad M_{\text{eff}} \approx \epsilon^{-1}, \quad N_{\text{eff}} \approx \epsilon^{-1} \log(\epsilon^{-2}). \quad (31)$$

Under (31), we see that

$$\lambda_1(1 - \gamma_M\lambda_1)^{2N_{\text{eff}}} \leq e^{-2N_{\text{eff}}\gamma_M} \lesssim \epsilon^2. \quad (32)$$

We now verify that $\text{ExcessRisk}(\mathbf{w}_{M+N}) \lesssim \epsilon$.

According to the proof of (18), it holds that

$$\frac{D_{\text{eff}}}{N_{\text{eff}}} \lesssim \gamma_M \text{tr}(\mathbf{H}) \lesssim \epsilon. \quad (33)$$

Now we choose $\mathbb{J} = \{1, 2\}$ so that

$$\begin{aligned} \frac{D_{\text{eff}}^{\text{finetune}}}{M_{\text{eff}}} &\lesssim \frac{1}{M_{\text{eff}}} \cdot \frac{\lambda_1(1 - \gamma_M\lambda_1)^{2N_{\text{eff}}}}{\mu_1} + \frac{1}{M_{\text{eff}}} \cdot \frac{\lambda_2(1 - \gamma_M\lambda_2)^{2N_{\text{eff}}}}{\mu_2} + 0 \\ &\lesssim \epsilon \cdot \frac{\epsilon^2}{\epsilon^2} + \epsilon \cdot \frac{\epsilon^{0.5} \cdot 1}{1} \quad (\text{by (32)}) \\ &\lesssim \epsilon \end{aligned} \quad (34)$$

For the bias error we have that

$$\begin{aligned} &\left\| \prod_{t=M}^{M+N-1} (\mathbf{I} - \gamma_t \mathbf{H}) \prod_{t=0}^{M-1} (\mathbf{I} - \gamma_t \mathbf{G})(\mathbf{w}_0 - \mathbf{w}^*) \right\|_{\mathbf{H}}^2 \\ &\leq \lambda_1(1 - \gamma_0\mu_1)^{2M_{\text{eff}}}(1 - \gamma_M\lambda_1)^{2N_{\text{eff}}}(\mathbf{w}^*[1])^2 + \lambda_2(1 - \gamma_0\mu_2)^{2M_{\text{eff}}}(1 - \gamma_M\lambda_2)^{2N_{\text{eff}}}(\mathbf{w}^*[2])^2 \\ &\leq (1 - \gamma_0\mu_1)^{2M_{\text{eff}}}(1 - \gamma_M\lambda_1)^{2N_{\text{eff}}} + \epsilon^{0.5}(1 - \gamma_0\mu_2)^{2M_{\text{eff}}}(1 - \gamma_M\lambda_2)^{2N_{\text{eff}}} \\ &\leq (1 - \gamma_M\lambda_1)^{2N_{\text{eff}}} + \epsilon^{0.5}(1 - \gamma_0\mu_2)^{2M_{\text{eff}}} \\ &\lesssim \epsilon. \quad (\text{by (32) and (31)}) \end{aligned} \quad (35)$$

(33), (34) and (35) together imply that $\text{ExcessRisk}(\mathbf{w}_{M+N}) \lesssim \epsilon$.

□