
In-Context Learning of a Linear Transformer Block: Benefits of the MLP Component and One-Step GD Initialization

Ruiqi Zhang
UC Berkeley
rqzhang@berkeley.edu

Jingfeng Wu
UC Berkeley
uuujf@berkeley.edu

Peter L. Bartlett
UC Berkeley and Google DeepMind
peter@berkeley.edu

Abstract

We study the *in-context learning* (ICL) ability of a *Linear Transformer Block* (LTB) that combines a linear attention component and a linear multi-layer perceptron (MLP) component. For ICL of linear regression with a Gaussian prior and a *non-zero mean*, we show that LTB can achieve nearly Bayes optimal ICL risk. In contrast, using only linear attention must incur an irreducible additive approximation error. Furthermore, we establish a correspondence between LTB and one-step gradient descent estimators with learnable initialization (GD- β), in the sense that every GD- β estimator can be implemented by an LTB estimator and every optimal LTB estimator that minimizes the in-class ICL risk is effectively a GD- β estimator. Finally, we show that GD- β estimators can be efficiently optimized with gradient flow, despite a non-convex training objective. Our results reveal that LTB achieves ICL by implementing GD- β , and they highlight the role of MLP layers in reducing approximation error.

1 Introduction

The recent dramatic progress in natural language processing can be attributed in large part to the development of large language models based on *Transformers* [34], such as BERT [11], LLaMA [32], PaLM [8], and the GPT series [28, 29, 6, 25]. In some pioneering pre-trained large language models, a new learning paradigm known as *in-context learning* (ICL) was observed (see, for example, [6, 7, 24, 20]). ICL refers to the capability of a pre-trained model to solve a new task based on a few in-context demonstrations without updating the model parameters.

A recent line of work quantifies ICL in tractable statistical learning setups such as linear regression with a Gaussian prior (see [14, 3] and references thereafter). Specifically, an ICL model takes a sequence $(\mathbf{x}_1, y_1, \dots, \mathbf{x}_M, y_M, \mathbf{x}_{\text{query}})$ as input and outputs a prediction of y_{query} , where $(\mathbf{x}_i, y_i)_{i=1}^M$ and $(\mathbf{x}_{\text{query}}, y_{\text{query}})$ are independent samples from an unknown task-specific distribution (where the task admits a prior distribution). The ICL model is pre-trained by fitting many empirical observations of such sequence-label pairs. Experiments show that Transformers can achieve an ICL risk close to that achieved by Bayes optimal estimators in many statistical learning tasks [14, 3].

For ICL of linear regression with a Gaussian prior, the data is generated as

$$\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{H}), \quad y \mid \tilde{\beta}, \mathbf{x} \sim \mathcal{N}(\tilde{\beta}^\top \mathbf{x}, \sigma^2),$$

where $\tilde{\beta}$ is a task parameter that satisfies a Gaussian prior, $\tilde{\beta} \sim \mathcal{N}(\mathbf{0}, \psi^2 \mathbf{I})$. In this setup, [14, 3] showed in experiments that Transformers can achieve nearly optimal ICL by matching the performance of the Bayes optimal estimator, that is, an optimally tuned ridge regression. Besides, [35] showed by construction that a single *linear self-attention* (LSA) can implement *one-step gradient descent with zero initialization* (GD-0), offering an insight into the ICL mechanism of the Transformer.

Later, theoretical works showed that optimal LSA models effectively correspond to GD-0 models [1], trained LSA models converge to GD-0 models [38], and GD-0 models (hence LSA models) can provably achieve nearly Bayes optimal ICL in certain regimes [37].

Our contributions. In this paper, we consider ICL in a more general setup, that is, linear regression with a Gaussian prior and a *non-zero mean* (that is, $\mathbb{E}[\tilde{\beta}] = \beta^* \neq 0$). The non-zero mean in task prior captures a common scenario where tasks share a signal. In this setting, we show that the LSA models considered in prior papers [1, 38, 37] incur an irreducible additive approximation error. Furthermore, we show this approximation error is mitigated by considering a *linear Transformer block* (LTB) that combines a linear multi-layer perceptron (MLP) component and an LSA component. Our results highlight the important role of MLP layers in Transformers in reducing the approximation error, and they suggest that theories about LSA [1, 38, 37] do not fully explain the power of Transformers. Motivated by this understanding, we investigate LTB in depth and obtain the following additional results:

- We show that LTB can implement *one-step gradient descent with learnable initialization*, referred to by us as GD- β . Additionally, we show that every optimal LTB estimator that minimizes the in-class ICL risk is *effectively* a GD- β estimator.
- Moreover, we show that the optimal GD- β , hence also the optimal LTB, nearly matches the performance of the Bayes optimal estimator for linear regression with a Gaussian prior and a non-zero mean, provided that the signal-to-noise ratio is upper bounded. These two results together suggest that LTB performs nearly optimal ICL by implementing GD- β .
- Finally, we show that the GD- β estimator can be efficiently optimized by gradient descent with an infinitesimal stepsize (that is, gradient flow) under the population ICL risk, despite the non-convexity of the objective.

Paper organization. The remaining paper is organized as follows. We conclude this section by introducing a set of notations. We then discuss related papers in Section 2. In Section 3, we set up our ICL problems and define LTB and LSA models mathematically. In Section 4, we show a positive approximation error gap between LSA and LTB, which highlights the importance of the MLP component and motivates our subsequent efforts to study LTB. In Section 5, we connect LSA estimators to GD- β estimators that are more interpretable for ICL of linear regression. In Section 6, we study the in-context learning and training of GD- β estimators. We conclude our paper and discuss future directions in Section 7.

Notations. We use lowercase bold letters to denote vectors and uppercase bold letters to denote matrices and tensors. For a vector $\mathbf{x}$ and a positive semi-definite (PSD) matrix $\mathbf{A}$, we write $\|\mathbf{x}\|_{\mathbf{A}} := \sqrt{\mathbf{x}^\top \mathbf{A} \mathbf{x}}$. We write $\langle \cdot, \cdot \rangle$ for the inner product, which is defined as $\langle \mathbf{x}, \mathbf{y} \rangle := \mathbf{x}^\top \mathbf{y}$ for vectors and $\langle \mathbf{A}, \mathbf{B} \rangle := \text{tr}(\mathbf{A} \mathbf{B}^\top)$ for matrices. We write $\mathbf{A}[i]$ as the i -th row of the matrix $\mathbf{A}$, $\mathbf{A}_{m,n}$ as the (m, n) -th entry, and $\mathbf{A}_{-1,-1}$ as the right-bottom entry. We write $\mathbf{0}_n, \mathbf{0}_{m \times n}, \mathbf{I}_n$ for the zero vector, zero matrix and identity matrix, respectively. We denote the Kronecker product as $\otimes$. For two matrices $\mathbf{A}$ and $\mathbf{B}$, $\mathbf{A} \otimes \mathbf{B}$ is a linear mapping which operates as $(\mathbf{A} \otimes \mathbf{B}) \circ \mathbf{C} = \mathbf{B} \mathbf{C} \mathbf{A}^\top$. For a positive semi-definite matrix $\mathbf{A}$, we write $\mathbf{A}^{\frac{1}{2}}$ as its principle square root, which is the unique positive semi-definite matrix $\mathbf{B}$ such that $\mathbf{B} \mathbf{B} = \mathbf{B} \mathbf{B}^\top = \mathbf{A}$. We also write $\mathbf{A}^{-\frac{1}{2}} = (\mathbf{A}^{\frac{1}{2}})^+$, where $(\cdot)^+$ denotes the Moore Penrose pseudo-inverse. For two sets A, B we write $A + B$ for the Minkowski sum, which is defined as $\{a + b : a \in A, b \in B\}$. We also write $a + B = \{a\} + B$ for an element a and a set B . We write $\text{null}(\cdot)$ for the null set of a matrix or a tensor.

2 Related Works

Empirical results for ICL in controlled settings. The work by [14] first considered ICL in controlled statistical learning setups. For noiseless linear regression, [14] showed in experiments that Transformers match the performance of the optimal estimator (that is, *Ordinary Least Square*). Subsequent works by [3, 18] extended their result to noisy linear regression with a Gaussian prior and showed that Transformers can match the performance of the Bayesian optimal estimator (that is, an optimally tuned ridge regression). Besides, [30] showed that the above holds even when Transformers are pretrained on a limited number of linear regression tasks. These papers only considered a Gaussian

prior with a zero mean. In contrast, we consider a Gaussian prior with a non-zero mean. Our setup better captures the common scenarios where the tasks share a signal.

The empirical investigation of ICL in tractable statistical learning setups goes beyond linear regression settings. For more examples, researchers empirically studied the ICL ability of Transformers for decision trees (see [14], they also considered two-layer networks), algorithm selection [4], linear mixture models [26], learning Fourier series [2], discrete boolean functions [5], representation learning [13, 16], and reinforcement learning [17, 19]. Among all these settings, Transformers can either compete with the Bayes optimal estimators or expert-designed strong benchmarks. These works are not directly comparable with our paper.

Transformer implements gradient descent. A line of work interpreted the ICL of Transformers by their abilities to implement *gradient descent* (GD) [35, 3, 10, 1, 38, 4, 37]. In experiments, [35, 10] showed that (multi-layer) Transformer outputs are close to (multi-step) GD outputs. When specialized to linear regression tasks, [35] constructed a single *linear self-attention* (LSA) that implements *one-step gradient descent with zero initialization* (GD-0). Subsequently, [1] showed that optimal LSA models effectively correspond to GD-0 models, [38] proved that trained LSA models converge to GD-0 models, [37] showed that GD-0 models (hence LSA models) can provably achieve nearly Bayes optimal ICL in certain regimes and provided a sharp task complexity analysis of the pre-training. These papers focused on LSA models. Instead, we consider a linear Transformer block that also utilizes the MLP layer.

From an approximation theory perspective, [3, 4] showed that Transformers can implement multi-step GD under general losses. In comparison, we consider a limited setting of linear regression and show the LSA models can implement one-step GD with learnable initialization (GD- β), moreover, every optimal LTB model is effectively an GD- β model. Both of our constructions utilize MLP layers in Transformers, highlighting its importance in reducing approximation error. Different from their results, we also show a negative approximation result that reveals the limitation of LSA models.

3 Preliminaries

Model input. We use $\mathbf{x} \in \mathbb{R}^d$ and $y \in \mathbb{R}$ to denote a feature vector and its label, respectively. Throughout the paper, we assume a fixed number of context examples, denoted by $M > 0$. We denote the context examples by $(\mathbf{X}, \mathbf{y}) \in \mathbb{R}^{M \times d} \times \mathbb{R}^M$, where each row represents a context example, denoted by $(\mathbf{x}_i^\top, y_i)$, $i = 1, \dots, M$. To formalize an ICL problem, the input of a model is a *token matrix* given by [1, 38]

$$\mathbf{E} := \begin{pmatrix} \mathbf{X}^\top & \mathbf{x} \\ \mathbf{y}^\top & 0 \end{pmatrix} \in \mathbb{R}^{(d+1) \times (M+1)}. \quad (3.1)$$

The output of a model corresponds to a prediction of y .

A Transformer block. Modern large language models are often made by stacking basic Transformer blocks (see, for example, [34]). A basic Transformer block consists of a *self-attention* layer and a *multi-layer perceptron* (MLP) layer [34], $\mathbf{E} \mapsto \text{MLP}[\text{ATTN}(\mathbf{E})]$. Here, the MLP layer is defined as $\text{MLP}(\mathbf{E}) := \mathbf{W}_2^\top \text{ReLU}(\mathbf{W}_1 \mathbf{E})$, where $\text{ReLU}(\cdot)$ refers to the entrywise *rectified linear unit* (ReLU) activation function, and $\mathbf{W}_1, \mathbf{W}_2 \in \mathbb{R}^{d_f \times (d+1)}$ are two weight matrices. The self-attention layer (we focus on the single-head version in this paper) is defined as $\text{ATTN}(\mathbf{E}) := \mathbf{E} + \mathbf{W}_P^\top \mathbf{W}_V \mathbf{E} \mathbf{M} \cdot \text{sfm}((\mathbf{W}_K \mathbf{E})^\top \mathbf{W}_Q \mathbf{E})$, where $\text{sfm}(\cdot)$ refers to the row-wise softmax operator, $\mathbf{W}_K, \mathbf{W}_Q \in \mathbb{R}^{d_k \times (d+1)}$ are the key and query matrices, respectively, $\mathbf{W}_P, \mathbf{W}_V \in \mathbb{R}^{d_v \times (d+1)}$ are the projection and value matrices, respectively, and $\mathbf{M}$ is a fixed masking matrix given by

$$\mathbf{M} := \begin{pmatrix} \mathbf{I}_M & 0 \\ 0 & 0 \end{pmatrix} \in \mathbb{R}^{(M+1) \times (M+1)}. \quad (3.2)$$

This mask matrix is included to reflect the asymmetric structure of a prompt since the label of query input $\mathbf{x}$ is not included in the token matrix [1, 21]. In the above formulation, d_k, d_v, d_f are three hyperparameters controlling the key size, value size, and width of the MLP layer, respectively. In the single-head case, it is common to set $d_k = d_v = d + 1$ and $d_f = 4(d + 1)$, where $d + 1$ corresponds to the embedding size (see, for example, [34]). Our formulation of a basic transformer block ignores all bias parameters and some popular techniques (such as layer normalization and dropout) to focus solely on the benefits brought by the model structure.

A linear Transformer block. To facilitate theoretical analysis, we ignore the non-linearities in the Transformer block (specifically, sfmx and ReLU) and work with a *linear Transformer block* (LTB) defined as

$$f_{\text{LTB}} : \mathbb{R}^{(d+1) \times (M+1)} \rightarrow \mathbb{R}, \quad \mathbf{E} \mapsto \left[\mathbf{W}_2^\top \mathbf{W}_1 \left(\mathbf{E} + \mathbf{W}_P^\top \mathbf{W}_V \mathbf{E} \mathbf{M} \frac{\mathbf{E}^\top \mathbf{W}_K^\top \mathbf{W}_Q \mathbf{E}}{M} \right) \right]_{-1, -1}, \quad (3.3)$$

where $\mathbf{E}$ is the token matrix given by (3.1), $[\cdot]_{-1, -1}$ refers to the bottom right entry of a matrix, $\mathbf{M}$ is the fixed masking matrix given by (3.2). Here, the trainable parameters are $\mathbf{W}_P, \mathbf{W}_V \in \mathbb{R}^{d_v \times (d+1)}$, $\mathbf{W}_K, \mathbf{W}_Q \in \mathbb{R}^{d_k \times (d+1)}$, and $\mathbf{W}_1, \mathbf{W}_2 \in \mathbb{R}^{d_f \times (d+1)}$. We use the bottom right entry of the transformed token matrix as the model output to form a prediction of the label y . The $1/M$ factor is a normalization factor in linear attention and can be absorbed into trainable parameters. Finally, we denote the hypothesis class formed by LTB models as $\mathcal{F}_{\text{LTB}} := \{f_{\text{LTB}} : \mathbf{W}_K, \mathbf{W}_Q, \mathbf{W}_V, \mathbf{W}_P, \mathbf{W}_1, \mathbf{W}_2, d_k \geq d, d_v \geq d+1, d_f \geq 1\}$, where f_{LTB} is defined in (3.3).

A linear self-attention. We will also consider a *linear self-attention* (LSA) defined as

$$f_{\text{LSA}} : \mathbb{R}^{(d+1) \times (M+1)} \rightarrow \mathbb{R}, \quad \mathbf{E} \mapsto \left[\mathbf{E} + \mathbf{W}_P^\top \mathbf{W}_V \mathbf{E} \mathbf{M} \frac{\mathbf{E}^\top \mathbf{W}_K^\top \mathbf{W}_Q \mathbf{E}}{M} \right]_{-1, -1}, \quad (3.4)$$

where $\mathbf{W}_K, \mathbf{W}_Q, \mathbf{W}_P, \mathbf{W}_V$ are trainable parameters. An LSA model can be viewed as an LTB model without the MLP layer (setting $d_f = d+1$ and $\mathbf{W}_1 = \mathbf{W}_2 = \mathbf{I}$). We remark that, unlike LTB, the residual connection in LSA plays no role because the bottom right entry of the prompt $\mathbf{E}$ is zero (see (3.1)). A variant of the LSA model has been studied by [1, 38, 37], where $\mathbf{W}_K^\top \mathbf{W}_Q$ and $\mathbf{W}_P^\top \mathbf{W}_V$ are respectively merged into one matrix parameter. Similarly, we denote the hypothesis class formed by LSA models as $\mathcal{F}_{\text{LSA}} := \{f_{\text{LSA}} : \mathbf{W}_K, \mathbf{W}_Q, \mathbf{W}_V, \mathbf{W}_P, d_k \geq d, d_v \geq 1\}$, where f_{LSA} is defined in (3.4).

Linear regression tasks with a shared signal. Assume that data and context examples are generated as follows.

Assumption 3.1 (Distributional conditions). Assume that $(\mathbf{X}, \mathbf{y}, \mathbf{x}, y)$ are generated by:

- First, a task parameter is independently generated by $\tilde{\beta} \sim \mathcal{N}(\beta^*, \Psi)$.
- The feature vectors are independently generated by $\mathbf{x}, \mathbf{x}_1, \dots, \mathbf{x}_M \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \mathbf{H})$.
- Then, the labels are generated by $y = \langle \tilde{\beta}, \mathbf{x} \rangle + \varepsilon$, $y_i = \langle \tilde{\beta}, \mathbf{x}_i \rangle + \varepsilon_i$, $i = 1, \dots, M$, where ε and ε_i 's are independently generated by $\varepsilon, \varepsilon_1, \dots, \varepsilon_M \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2)$.

Here, $\sigma^2 \geq 0$, $\mathbf{H} \succeq \mathbf{0}$, $\Psi \succeq \mathbf{0}$, and $\beta^* \in \mathbb{R}^d$ are fixed but unknown quantities that govern the data distribution. We denote $\varepsilon = (\varepsilon_1, \dots, \varepsilon_M)^\top$.

We emphasize the importance of the mean of the task parameter β^* in Assumption 3.1. A non-zero β^* represents a shared signal across tasks, which is arguably common in practice. This assumption is implicitly used in [12] where they assumed task parameters are close to a meta parameter. In comparison, the prior works for ICL of linear regression [1, 38, 37] only considered a special case where $\beta^* = 0$. In this special case, they showed that a single LSA layer can achieve nearly optimal ICL by approximating one-step gradient descent from zero initialization (GD-0). In more general cases where β^* is non-zero, we will show in Section 4 that LSA is insufficient to learn the shared signal and must incur an irreducible approximation error compared to LTB models. This sets our results apart from the prior papers.

ICL risk. We measure the ICL risk of a model f by the mean squared error,

$$\mathcal{R}(f) := \mathbb{E}(f(\mathbf{E}) - y)^2, \quad (3.5)$$

where $\mathbf{E}$ is defined in (3.1) and the expectation is over $\mathbf{E}$ (equivalent to over $\mathbf{X}, \mathbf{y}$, and $\mathbf{x}$) and y .

4 Benefits of the MLP Component

Our first main result separates the approximation abilities of the LTB and LSA models for ICL. Recall that an LSA model can be viewed as a special case of the LTB model with $\mathbf{W}_2 = \mathbf{W}_1 = \mathbf{I}$. So the

best ICL risk achieved by LTB models is no larger than the best ICL risk achieved by LSA models. However, our next theorem shows a strictly positive gap between the best ICL risks achieved by those two model classes. This result highlights the benefits of the MLP layer for reducing approximation error in Transformer.

Theorem 4.1 (Approximation gap). *Consider the ICL risk defined by (3.5) and the two hypothesis classes $\mathcal{F}_{\text{LTB}}$ and $\mathcal{F}_{\text{LSA}}$. Suppose that Assumption 3.1 holds. Then we have*

- $\inf_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f)$ is independent of β^* .
- $\inf_{f \in \mathcal{F}_{\text{LSA}}} \mathcal{R}(f)$ is a function of β^* . Moreover,

$$\inf_{f \in \mathcal{F}_{\text{LSA}}} \mathcal{R}(f) - \inf_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f) \geq \max \left\{ \frac{2}{3(M+1)}, \frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2}{(M+1)^2 \text{tr}((\mathbf{H}\Psi)^2)} \right\} \|\beta^*\|_{\mathbf{H}}^2.$$

The proof of Theorem 4.1 is deferred to Appendix B. Theorem 4.1 reveals a gap in terms of the approximation abilities between the hypothesis set of LTB models and that of LSA models. Specifically, the best ICL performance achieved by LTB models is independent of β^* , while the best ICL performance achieved by LSA models is sensitive to the norm of β^* . In particular, when $\|\beta^*\|_{\mathbf{H}}^2 = \Omega(M)$, the best ICL risk of the former is smaller than that of the latter by at least a *constant* additive term. So when β^* is large, the hypothesis class formed by LSA models is restricted in its ability to perform effective ICL. We will show in Section 5 that the hypothesis class formed by LTB models can achieve nearly optimal ICL in this case.

We also remark that the $\Theta(1/M)$ factor in the lower bound in Theorem 4.1 is not improvable. This is because LSA models can implement a one-step GD algorithm that is *consistent* for linear regression tasks (that is, the risk converges to the Bayes risk as the number of context examples goes to infinity), with an excess risk bound of $\Theta(1/M)$ [38, 37]. So the approximation error gap is at most $\Theta(1/M)$. Nonetheless, the $\Omega(\|\beta^*\|_{\mathbf{H}}^2/M)$ approximation gap between LSA models and LTB models shown by Theorem 4.1 suggests that β^* is not learnable by LSA models during pre-training. In contrast, we will show in Section 6 that LTB models can learn β^* during pre-training.

We emphasize that the ability of LTB to learn non-zero mean is a joint effect of an MLP component and a skip connection. Note that LSA also has a skip connection, but its skip connection is inactive. In comparison, the MLP component in LTB activates the skip connection. Therefore, we attribute the ability to learn non-zero mean to the MLP component, which is the only difference between LTB and LSA. Nonetheless, one can attribute the ability to learn non-zero mean to the skip connection: without a skip connection, LTB reduces to LSA with a potential rank constraint on the parameter, which cannot learn non-zero mean as we have proved. The above two explanations take different perspectives to interpret the same phenomenon. Finally, we remark that Theorem 4.1 holds even when $\mathbf{H}$ and Ψ in Assumption 3.1 are not full rank.

Does scratchpad help? We have demonstrated that employing a single-layer LSA introduces an additional approximation error in the in-context learning problem for the linear regression task defined in Assumption 3.1. Nonetheless, are there alternative structures, apart from MLPs, that could potentially reduce this approximation error? One plausible strategy is to include a “scratchpad” in the input token. Specifically, we construct the token matrix as follows:

$$\mathbf{E} := \begin{pmatrix} \mathbf{X}^\top & \mathbf{x} \\ \mathbf{1}_M^\top & 1 \\ \mathbf{y}^\top & 0 \end{pmatrix} \in \mathbb{R}^{(d+2) \times (M+1)} \quad (4.1)$$

which we then input into the LSA layer. We discuss the limitation of this scheme in Appendix J. Notably, this method does not successfully recover the GD- β estimator defined in Section 5. We leave it as future work to see whether the token matrix with scratchpad could implement other types of estimators that more effectively address the linear regression tasks defined in Assumption 3.1, as well as whether additional structures could help alleviate this approximation error.

Experiments on GPT2. Theorem 4.1 shows the importance of the MLP component in LSA models for reducing approximation error. We also empirically validate this result by training a more complex GPT2 model [14] for the ICL tasks specified by Assumption 3.1. In the experiments, we use a GPT2-small model (with or without the MLP component) with 6 layers and 4 heads in each layer. The experiments follow the setting in [14], except that we train the model using a token matrix defined

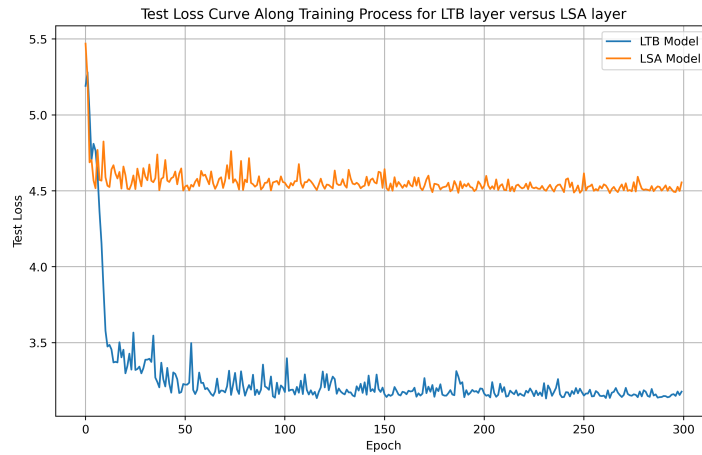


Figure 1: The test loss along the training process for LTB and LSA layer.

in (3.1). We consider two ICL settings, which instantiates Assumption 3.1 with $\beta^* = (0, 0, \dots, 0)^\top$ and $\beta^* = (10, 10, \dots, 10)^\top$, respectively. We set $d = 20, M = 40, \sigma = 0, \Psi = \mathbf{H} = \mathbf{I}_d$ in both settings. More experimental details are in Appendix I. In each setting, we train and test the model using the same data distribution. The experimental results are presented in Table 1. From Table 1, we observe that both models (with or without the MLP component) achieve a nearly zero loss when the task mean is zero. However, when the task mean is set away from zero, the GPT2 model with MLP component still performs relatively well while the GPT2 model without MLP component incurs a significantly larger loss. These empirical observations are consistent with our Theorem 4.1, indicating the benefits of MLP layers in reducing approximation error for ICL of linear regression with a shared signal.

Experiments on LSA and LTB. We trained both the LSA and LTB layers for $\beta^* = (1, 1, \dots, 1)^\top$, and found that the trained LTB layer consistently achieved a significantly lower ICL risk than the LSA layer (see Figure 1). In these experiments, we adhered strictly to the previously defined LSA and LTB structures, setting the parameters as follows: $d = 5, M = 5, \sigma = 0$, and $\Psi = \mathbf{H} = \mathbf{I}_d$. At each training step, we sampled $B = 128$ new linear regression tasks.

Table 1: Losses of GPT2 with or without MLP component for linear regression with a shared signal.

Model	GPT2	GPT2-noMLP
$\beta^* = 0 \times \mathbf{1}$	0.003	0.024
$\beta^* = 10 \times \mathbf{1}$	0.013	8.871

In this part, we have shown that LSA models, the primary focus in previous ICL theory literature (see, e.g., [1, 38, 37] and references therein), are not sufficiently expressive for ICL of linear regression with a shared signal. In what follows, we will show LTB models are sufficient for this ICL problem.

5 LTB Implements One-Step GD with Learnable Initialization

To understand the expressive power of the LTB models, we build a connection between $\mathcal{F}_{\text{LTB}}$ and its subset of models which we call *one-step GD with learnable initialization* (GD- β). We will first introduce GD- β models and then show that the best LTB models that minimize the ICL risk effectively belong to GD- β models.

The GD- β models. A GD- β model is defined as

$$f_{\text{GD-}\beta} : \mathbb{R}^{(d+1) \times (M+1)} \rightarrow \mathbb{R}, \quad \mathbf{E} \mapsto \left\langle \beta - \Gamma \cdot \frac{1}{M} \mathbf{X}^\top (\mathbf{X}\beta - \mathbf{y}), \mathbf{x} \right\rangle, \quad (5.1)$$

where $\mathbf{E}$ is the token matrix given by (3.1), $\Gamma \in \mathbb{R}^{d \times d}$ and $\beta \in \mathbb{R}^d$ are trainable parameters. Similarly, we define the function class formed by GD- β models as $\mathcal{F}_{\text{GD-}\beta} := \{f_{\text{GD-}\beta} : \beta \in \mathbb{R}^d, \Gamma \in \mathbb{R}^{d \times d}\}$.

A GD- β model computes a parameter that fits the context examples $(\mathbf{X}, \mathbf{y})$ and uses that parameter to make a linear prediction of label y on feature $\mathbf{x}$. More specifically, the first step is by using one gradient descent step with a *matrix stepsize* $\mathbf{\Gamma}$ and an *initialization* β on the least square objective formed by context examples $(\mathbf{X}, \mathbf{y})$.

LTB implements GD- β . A GD- β model (5.1) is a special case of an LTb model (3.3) by setting

$$\mathbf{W}_2^\top \mathbf{W}_1 = \begin{pmatrix} * & * \\ \beta^\top & 1 \end{pmatrix}, \quad \mathbf{W}_P^\top \mathbf{W}_V = \begin{pmatrix} -\mathbf{I}_d & \mathbf{0}_d \\ \mathbf{0}_d^\top & 1 \end{pmatrix}, \quad \mathbf{W}_K^\top \mathbf{W}_Q = \begin{pmatrix} \mathbf{\Gamma} & * \\ \mathbf{0}_d^\top & * \end{pmatrix},$$

where $*$ denotes the entries that do not affect the model output (hence can be set to anything). Note that $\mathbf{\Gamma}$ and β are *free* provided that $d_v \geq d + 1$, $d_k \geq d$ and $d_f \geq 1$ (as required in $\mathcal{F}_{\text{LTb}}$). In sum, we have proved the following lemma showing that the set of GD- β models belongs to the set of LTb models.

Lemma 5.1. *We have $\mathcal{F}_{\text{GD-}\beta} \subseteq \mathcal{F}_{\text{LTb}}$. Therefore, $\inf_{f \in \mathcal{F}_{\text{GD-}\beta}} \mathcal{R}(f) \geq \inf_{f \in \mathcal{F}_{\text{LTb}}} \mathcal{R}(f)$.*

Optimal GD- β models for ICL. We now consider $\mathcal{F}_{\text{GD-}\beta}$ and its optimal ICL risk. We have the following theorem that computes the globally minimal ICL risk over $\mathcal{F}_{\text{GD-}\beta}$ and specifies the sufficient and necessary conditions for a global minimizer. The proof is deferred to Appendix C.

Theorem 5.2 (Optimal GD- β models). *Consider the ICL risk defined by (3.5). Suppose that Assumption 3.1 holds. Then we have*

- The minimal ICL risk of $\mathcal{F}_{\text{GD-}\beta}$ is

$$\inf_{f \in \mathcal{F}_{\text{GD-}\beta}} \mathcal{R}(f) = \sigma^2 + \text{tr} \left(\mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \left(\mathbf{I} - \mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \mathbf{\Omega}^{-1} \right) \right), \quad (5.2)$$

where $\mathbf{\Omega} := [(M + 1) \mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} + (\text{tr}(\mathbf{H} \mathbf{\Psi}) + \sigma^2) \mathbf{I}_d] / M$.

- The global optimal parameters for $\mathcal{F}_{\text{GD-}\beta}$ that attain the minimum (5.2) take the following form:

$$\beta = \beta^* + \text{null}(\mathbf{H}), \quad \mathbf{\Gamma} = \mathbf{\Gamma}^* + \text{null}(\mathbf{H}^{\otimes 2}), \quad \text{where} \quad \mathbf{\Gamma}^* := \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \mathbf{\Omega}^{-1} \mathbf{H}^{-\frac{1}{2}} \quad (5.3)$$

and β^* is the mean of the task parameter in Assumption 3.1 Here, $\text{null}(\mathbf{H}) := \{ \mathbf{z} \in \mathbb{R}^d : \mathbf{H} \mathbf{z} = \mathbf{0} \}$ is the null space of $\mathbf{H}$, and $\text{null}(\mathbf{H}^{\otimes 2}) = \{ \mathbf{Z} \in \mathbb{R}^{d \times d} : \mathbf{H} \mathbf{Z} \mathbf{H} = \mathbf{0} \}$ is the null space of $\mathbf{H}^{\otimes 2}$. In particular, when $\mathbf{H}$ is positive definite, the global optimal parameter is unique, $(\beta, \mathbf{\Gamma}) = (\beta^*, \mathbf{\Gamma}^*)$.

- Under Assumption 3.1, the global optimal $f_{\text{GD-}\beta}$ that attains the minimum (5.2) is unique as a function of $\mathbf{E}$ (given by (3.1)) and takes the following form:

$$f^*(\mathbf{E}) = \left\langle \beta^* - \frac{1}{M} \mathbf{\Gamma}^* \mathbf{X}^\top (\mathbf{X} \beta^* - \mathbf{y}), \mathbf{x} \right\rangle. \quad (5.4)$$

Theorem 5.2 characterizes the optimal GD- β models for ICL. In the above theorem, $\mathbf{\Gamma}^* \rightarrow \mathbf{H}^{-1}$ as the context length M goes to infinity (assuming that $\mathbf{H}$ is positive definite). In this case, the optimal GD- β function (5.2) implements one Newton step from initialization β^* . With a finite context length M , (5.2) implements one regularized Newton step from initialization β^* .

Optimal LTb models for ICL. Lemma 5.1 shows that the best ICL risk achieved by an GD- β model is no smaller than the best ICL risk achieved by an LTb model. Surprisingly, our next theorem shows that the best ICL risk achieved by a GD- β is *equal* to that achieved by an LTb model. Therefore, the hypothesis set $\mathcal{F}_{\text{GD-}\beta}$ is diverse enough to match the approximation ability of the larger hypothesis set $\mathcal{F}_{\text{LTb}}$ for ICL.

Theorem 5.3 (Optimal LTb models). *Consider the ICL risk defined by (3.5). Suppose that Assumption 3.1 holds and that $\text{rank}(\mathbf{H}^{\frac{1}{2}} \mathbf{\Psi}^{\frac{1}{2}}) \geq 2$. Then we have*

- The minimal ICL risk of $\mathcal{F}_{\text{LTb}}$ and of $\mathcal{F}_{\text{GD-}\beta}$ are equal,

$$\inf_{f \in \mathcal{F}_{\text{LTb}}} \mathcal{R}(f) = \inf_{f \in \mathcal{F}_{\text{GD-}\beta}} \mathcal{R}(f) = \text{RHS of (5.2)}$$

- Rewrite an LTB model as f_{LTB} in (3.3) with parameters ($*$ denotes parameters that do not affect the output)

$$\mathbf{W}_2^\top \mathbf{W}_1 = \begin{pmatrix} * & * \\ \gamma^\top & * \end{pmatrix}, \quad \mathbf{W}_K^\top \mathbf{W}_Q = \begin{pmatrix} \mathbf{V}_{11} & * \\ \mathbf{v}_{12}^\top & * \end{pmatrix}, \quad \mathbf{W}_2^\top \mathbf{W}_1 \mathbf{W}_P^\top \mathbf{W}_V = \begin{pmatrix} * & * \\ \mathbf{v}_{21}^\top & v_{-1} \end{pmatrix}.$$

Then the sufficient and necessary conditions for $f_{\text{LTB}} \in \arg \min_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f)$ are

$$v_{-1} \neq 0, \quad v_{-1} \mathbf{v}_{12} \in \text{null}(\mathbf{H}), \quad \mathbf{v}_{21} \in -v_{-1} \beta^* + \text{null}(\mathbf{H}), \quad \gamma \in \beta^* + \text{null}(\mathbf{H}), \\ v_{-1} \mathbf{V}_{11} \in \mathbf{\Gamma}^{*\top} - v_{-1} \beta^* \mathbf{v}_{12}^\top + \text{null}(\mathbf{H}^{\otimes 2}),$$

where β^* is defined in Assumption 3.1 and $\mathbf{\Gamma}^*$ is defined in (5.3). In particular, when $\mathbf{H}$ is positive definite, the globally optimal parameter represented by $(\mathbf{V}_{11}, \mathbf{v}_{12}, \mathbf{v}_{21}, v_{-1}, \gamma)$ is unique up to a rescaling of v_{-1} .

- Under Assumption 3.1, the globally optimal LTB model (that is, a function in $\arg \min_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f)$) is unique as a function of $\mathbf{E}$ (given by (3.1)) and takes the form of (5.4) almost surely.

Lemma 5.1 and Theorem 5.3 together show that $\mathcal{F}_{\text{GD-}\beta}$ is a representative subset of $\mathcal{F}_{\text{LTB}}$ that does not incur additional approximation error. In addition, every optimal LTB model is effectively an optimal GD- β model when restricted to all possible token matrices. Note that the optimal model parameters for LTB or GD- β are not unique because of redundant parameterization. But the optimal LTB and GD- β models are unique as a function of all possible token matrices. The above holds even when $\mathbf{H}$ and Ψ are potentially rank deficient.

Comparison with prior works. A line of papers considers LSA models for ICL under Assumption 3.1 with $\beta^* = 0$ [35, 1, 38, 37]. They show that LSA models can (effectively) implement all possible GD-0 models that specialize GD- β models by fixing $\beta = \mathbf{0}$. In addition, they show that every optimal LSA model is (effectively) a GD-0 model for ICL under Assumption 3.1 with $\beta^* = \mathbf{0}$ (see, for example, Theorem 1 in [1]). In comparison, we consider a harder ICL problem that allows a large shared signal in tasks (that is, a large β^* in Assumption 3.1). In this setting, our Theorem 4.1 shows that $\mathcal{F}_{\text{LSA}}$ (hence its subset formed by GD-0 models), as a subset of $\mathcal{F}_{\text{LTB}}$, incurs an additional approximation error proportional to $\|\beta^*\|_{\mathbf{H}}^2$ compared with $\mathcal{F}_{\text{LTB}}$. In contrast, $\mathcal{F}_{\text{GD-}\beta}$, as a subset of $\mathcal{F}_{\text{LTB}}$, does not incur additional approximation error according to our Theorem 5.3. Thus the LSA and GD-0 models considered by [35, 1, 38, 37] are not capable of learning the shared signal β^* , while an LTB model can learn β^* through implementing GD- β and encoding β^* in the initialization parameter.

6 Training and In-Context Learning of GD-beta

We have shown that $\mathcal{F}_{\text{GD-}\beta}$ is a representative subset of $\mathcal{F}_{\text{LTB}}$ that effectively contains every optimal LTB model. We now examine the ICL and training of GD- β models.

Nearly optimal ICL with GD- β . We will compare the best ICL risk achieved by GD- β with the best ICL risk achieved by any estimator. The following lemma is an extension of Proposition 5.1 and Corollary 5.2 in [37] (which is based on [33]) that characterizes the Bayes optimal ICL risk among all estimators.

Lemma 6.1 (Bayes optimal ICL). *Given a task-specific dataset $(\mathbf{X}, \mathbf{y}, \mathbf{x}, y)$ sampled according to Assumption 3.1, let $g(\mathbf{X}, \mathbf{y}, \mathbf{x})$ be an arbitrary estimator for y and measure the average linear regression risk by $\mathcal{L}(g; \mathbf{X}) := \mathbb{E}[(g(\mathbf{X}, \mathbf{y}, \mathbf{x}) - y)^2 | \mathbf{X}]$. It is clear that $\mathbb{E}\mathcal{L}(g; \mathbf{X}) = \mathcal{R}(g)$. Then,*

- The optimal estimator that minimizes the average linear regression risk $\mathcal{L}(\cdot; \mathbf{X})$ is $g^*(\mathbf{X}, \mathbf{y}, \mathbf{x}) = \mathbf{x}^\top \beta^* + \mathbf{x}^\top \Psi^{\frac{1}{2}} (\Psi^{\frac{1}{2}} \mathbf{X}^\top \mathbf{X} \Psi^{\frac{1}{2}} + \sigma^2 \mathbf{I}_d)^{-1} \Psi^{\frac{1}{2}} \mathbf{X}^\top (\mathbf{y} - \mathbf{X} \beta^*)$.
- Assume the signal-to-noise ratio is upper bounded, that is, $\text{tr}(\mathbf{H} \Psi) \lesssim \sigma^2$, then with probability at least $1 - \exp(-\Omega(M))$ over the randomness of $\mathbf{X}$, it holds that $\mathcal{L}(g^*; \mathbf{X}) - \sigma^2 \simeq \sum_{i=1}^d \min\{\bar{\phi}, \phi_i\}$, where $\bar{\phi} \approx \frac{\sigma^2}{M}$, and $(\phi_i)_{i \geq 1}$ are the eigenvalues of $\Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}}$.

The proof is deferred to Appendix F. Lemma 6.1 shows that the Bayes optimal estimator is a ridge regression estimator centered at β^* . This is consistent with [37] where the Bayes optimal estimator

is a ridge regression estimator as they assumed $\beta^* = 0$. The following corollary of Theorem 5.2 computes the rate of the ICL risk achieved by the optimal GD- β model.

Corollary 6.2. *Under the setup of Theorem 5.2, additionally assume the signal-to-noise ratio is upper bounded, that is, $\text{tr}(\Psi\mathbf{H}) \lesssim \sigma^2$. Then we have $\inf_{f \in \mathcal{F}_{\text{GD-}\beta}} \mathcal{R}(f) - \sigma^2 \simeq \sum_{i=1}^d \min\{\bar{\phi}, \phi_i\}$, where $(\phi_i)_{i \geq 0}$ are the eigenvalues of $\Psi^{\frac{1}{2}}\mathbf{H}\Psi^{\frac{1}{2}}$ and $\bar{\phi} := [\text{tr}(\Psi^{\frac{1}{2}}\mathbf{H}\Psi^{\frac{1}{2}}) + \sigma^2]/M \simeq \sigma^2/M$.*

The optimal (expected) ICL risk achieved by $\mathcal{F}_{\text{GD-}\beta}$ in Corollary 6.2 matches the (high probability) Bayes optimal ICL risk in Lemma 6.1 ignoring constant factors, provided that the signal-to-noise ratio is upper bounded. Therefore $\mathcal{F}_{\text{GD-}\beta}$ achieves nearly Bayes optimal ICL risk. As a consequence, the larger hypothesis set $\mathcal{F}_{\text{LTB}}$ also achieves nearly Bayes optimal ICL of linear regression under Assumption 3.1.

For the simplicity of discussion, we assume a fixed context length during pretraining and inference. Our discussions can be extended to allow a different context length during pretraining and inference using techniques in [37]. However, this is not the main focus of this work.

Optimization of GD- β with infinite tasks. We have shown that $\mathcal{F}_{\text{GD-}\beta}$ is a representative subset of $\mathcal{F}_{\text{LTB}}$ that covers the optimal LTB models and achieves nearly optimal ICL risk. We now consider the optimization in the parameter space specified by $\mathcal{F}_{\text{GD-}\beta}$. For simplicity, we follow [38] and consider gradient descent with an infinitesimal stepsize on the ICL objective with an infinite number of tasks. That is, we consider the optimization of gradient flow on the population ICL risk under the parameterization of GD- β ,

$$\frac{d\beta(t)}{dt} = -\frac{\partial}{\partial\beta}\mathcal{R}(f_{\text{GD-}\beta}), \quad \frac{d\Gamma(t)}{dt} = -\frac{\partial}{\partial\Gamma}\mathcal{R}(f_{\text{GD-}\beta}), \quad (6.1)$$

where $\mathcal{R}$ is defined by (3.5) and $f_{\text{GD-}\beta}$ is defined by (5.1).

The following theorem guarantees the global convergence of gradient flow. We introduce some notation to accommodate cases when $\mathbf{H}$ is rank deficient. Let $\mathcal{P}_{\mathcal{S}}$ be the orthogonal projection operator onto a subspace $\mathcal{S}$. Let $\mathcal{H} = \text{Im}(\mathbf{H})$ be the image space of matrix $\mathbf{H}$ (viewing $\mathbf{H}$ as a linear map) and $\mathcal{H}^{\perp} := \text{null}(\mathbf{H})$ be its orthogonal complement. Similarly, let $\mathcal{Z} := \text{Im}(\mathbf{H}^{\otimes 2}) = \{\mathbf{H}\mathbf{Z}\mathbf{H}, \mathbf{Z} \in \mathbb{R}^{d \times d}\}$ be the image space of the operator $\mathbf{H}^{\otimes 2}$ and $\mathcal{Z}^{\perp}$ be its orthogonal complement. Then we have the following theorem.

Theorem 6.3. *Consider the gradient flow defined by (6.1) with initialization $\beta(0), \Gamma(0)$. We have,*

$$\begin{aligned} \mathcal{P}_{\mathcal{H}}(\beta(t)) &\rightarrow \mathcal{P}_{\mathcal{H}}(\beta^*), & \mathcal{P}_{\mathcal{H}^{\perp}}(\beta(t)) &= \mathcal{P}_{\mathcal{H}^{\perp}}(\beta(0)), \\ \mathcal{P}_{\mathcal{Z}}(\Gamma(t)) &\rightarrow \mathcal{P}_{\mathcal{Z}}(\Gamma^*), & \mathcal{P}_{\mathcal{Z}^{\perp}}(\Gamma(t)) &= \mathcal{P}_{\mathcal{Z}^{\perp}}(\Gamma(0)) \end{aligned}$$

as $t \rightarrow \infty$. In particular, if $\mathbf{H}$ is positive definite, the gradient flow converges to the unique global minimizer of ICL risk over GD- β class, that is, $\beta(t) \rightarrow \beta^*$ and $\Gamma(t) \rightarrow \Gamma^*$ as $t \rightarrow \infty$.

The proof, as well as the convergence rate, is deferred to Appendix G. We remark that (6.1) is a complex dynamical system on a non-convex potential function of β and Γ . We briefly discuss our proof techniques assuming that $\mathbf{H}$ is full rank. The rank-deficient cases can be handled in the same way by applying appropriate project operators. To conquer the non-convex optimization issue, we observe that for every fixed Γ , the potential as a function of β is smooth and strongly convex with a uniformly bounded condition number. This observation allows us to establish a uniform convergence for β . When β is sufficiently close to β^* , the potential as a function of Γ is approximately convex, allowing us to track the convergence of Γ .

Theorem 6.3 shows that optimization of GD- β can be done efficiently by gradient flow without suffering from non-convexity. However, as we have shown in previous sections, LTB utilizes a more complex parameterization than GD- β . So Theorem 6.3 does not imply optimization of LTB is easy. We leave it as future work to study the optimization and statistical complexity for directly learning LTB models.

7 Concluding Remarks

In this paper, we study the in-context learning of linear regression with a shared signal represented by a Gaussian prior with a non-zero mean. We show that although the linear self-attention layer

discussed in prior works is consistent for this more complex task, its risk has an inevitable gap compared to that of the linear Transformer block (LTB), which is a linear self-attention layer followed by a linear multi-layer perception (MLP) layer. Next, we show that the effectiveness of the LTB arises because it can implement the one-step gradient descent estimator with learnable initialization (GD- β). Moreover, all global minimizers in the LTB class are equivalent to the unique global minimizer in the GD- β class, which can achieve nearly Bayes optimal in-context learning risk. Finally, we consider training on in-context examples and prove global convergence over the GD- β class of gradient flow on the population loss. Several future directions are worth discussing.

Optimization and statistical complexity. This paper provides an approximation theory of LTB and an optimization theory of GD- β . However, the statistical complexity of learning LTB or GD- β is not considered. The work by [37] provided techniques for analyzing the statistical task complexity for pre-training GD-0. An interesting direction is to extend their method to study the statistical complexity of learning GD- β . However, their method crucially relies on the convexity of the risk induced GD-0, while we have shown that the risk induced by GD- β is non-convex. New ideas for dealing with non-convexity are needed here.

From LTB to Transformer block. We focus on LTB in this work, which simplifies a vanilla Transformer block by removing the non-linearities from the softmax self-attention and the ReLU activation in the MLP layers. This simplification allows us to obtain precise theoretical results for LTB (such as its connection to GD- β). On the other hand, non-linearities are arguably necessary for Transformers to work well in practice. An important next step is to further consider the theoretical benefits of non-linearities based on our current results.

Roles of MLP layers. The work by [15] (and references thereafter) empirically found that MLP layers operate as key-value memories that store human-interpretable patterns in some pre-trained Transformers. Their work motivated a method for locating and editing information stored in language models by modifying their MLP layers (see, e.g., [9, 23]). Our work proves that the MLP component enables LTB to learn the shared signal in linear regression tasks, which cannot be done by a single LSA component. We leave it as future work to theoretically clarify the information stored in the MLP component.

Acknowledgments

We gratefully acknowledge the support of the NSF and the Simons Foundation for the Collaboration on the Theoretical Foundations of Deep Learning through awards DMS-2031883 and #814639, and of the NSF through grant DMS-2023505.

References

- [1] Kwangjun Ahn, Xiang Cheng, Hadi Daneshmand, and Suvrit Sra. Transformers learn to implement preconditioned gradient descent for in-context learning. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- [2] Kabir Ahuja, Madhur Panwar, and Navin Goyal. In-context learning through the bayesian prism. In *The Twelfth International Conference on Learning Representations*, 2024.
- [3] Ekin Akyürek, Dale Schuurmans, Jacob Andreas, Tengyu Ma, and Denny Zhou. What learning algorithm is in-context learning? investigations with linear models. In *The Eleventh International Conference on Learning Representations*, 2022.
- [4] Yu Bai, Fan Chen, Huan Wang, Caiming Xiong, and Song Mei. Transformers as statisticians: Provable in-context learning with in-context algorithm selection. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- [5] Satwik Bhattamishra, Arkil Patel, Phil Blunsom, and Varun Kanade. Understanding in-context learning in transformers and LLMs by learning to learn discrete functions. In *The Twelfth International Conference on Learning Representations*, 2024.
- [6] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901, 2020.

- [7] Yanda Chen, Ruiqi Zhong, Sheng Zha, George Karypis, and He He. Meta-learning via language model in-context tuning. *arXiv preprint arXiv:2110.07814*, 2021.
- [8] Aakanksha Chowdhery, Sharan Narang, Jacob Devlin, Maarten Bosma, Gaurav Mishra, Adam Roberts, Paul Barham, Hyung Won Chung, Charles Sutton, Sebastian Gehrmann, et al. Palm: Scaling language modeling with pathways. *arXiv preprint arXiv:2204.02311*, 2022.
- [9] Damai Dai, Li Dong, Yaru Hao, Zhifang Sui, Baobao Chang, and Furu Wei. Knowledge neurons in pretrained transformers. In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 8493–8502, 2022.
- [10] Damai Dai, Yutao Sun, Li Dong, Yaru Hao, Shuming Ma, Zhifang Sui, and Furu Wei. Why can GPT learn in-context? language models secretly perform gradient descent as meta-optimizers. In *Findings of the Association for Computational Linguistics: ACL 2023*, 2023.
- [11] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*, 2018.
- [12] Chelsea Finn, Pieter Abbeel, and Sergey Levine. Model-agnostic meta-learning for fast adaptation of deep networks. In *International conference on machine learning*, pages 1126–1135. PMLR, 2017.
- [13] Jingwen Fu, Tao Yang, Yuwang Wang, Yan Lu, and Nanning Zheng. How does representation impact in-context learning: A exploration on a synthetic task. *arXiv preprint arXiv:2309.06054*, 2023.
- [14] Shivam Garg, Dimitris Tsipras, Percy S Liang, and Gregory Valiant. What can transformers learn in-context? a case study of simple function classes. *Advances in Neural Information Processing Systems*, 35:30583–30598, 2022.
- [15] Mor Geva, Roei Schuster, Jonathan Berant, and Omer Levy. Transformer feed-forward layers are key-value memories. *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, 2021.
- [16] Tianyu Guo, Wei Hu, Song Mei, Huan Wang, Caiming Xiong, Silvio Savarese, and Yu Bai. How do transformers learn in-context beyond simple functions? a case study on learning with representations. In *The Twelfth International Conference on Learning Representations*, 2024.
- [17] Jonathan N Lee, Annie Xie, Aldo Pacchiano, Yash Chandak, Chelsea Finn, Ofir Nachum, and Emma Brunskill. Supervised pretraining can learn in-context reinforcement learning. *arXiv preprint arXiv:2306.14892*, 2023.
- [18] Yingcong Li, Muhammed Emrullah Ildiz, Dimitris Papailiopoulos, and Samet Oymak. Transformers as algorithms: Generalization and stability in in-context learning. In *International Conference on Machine Learning*, pages 19565–19594. PMLR, 2023.
- [19] Licong Lin, Yu Bai, and Song Mei. Transformers as decision makers: Provable in-context reinforcement learning via supervised pretraining. In *The Twelfth International Conference on Learning Representations*, 2024.
- [20] Pengfei Liu, Weizhe Yuan, Jinlan Fu, Zhengbao Jiang, Hiroaki Hayashi, and Graham Neubig. Pre-train, prompt, and predict: A systematic survey of prompting methods in natural language processing. *ACM Computing Surveys*, 55(9):1–35, 2023.
- [21] Arvind Mahankali, Tatsunori B Hashimoto, and Tengyu Ma. One step of gradient descent is provably the optimal in-context learner with one layer of linear self-attention. In *The Twelfth International Conference on Learning Representations*, 2024.
- [22] AR Meenakshi and C Rajian. On a product of positive semidefinite matrices. *Linear algebra and its applications*, 295(1-3):3–6, 1999.
- [23] Kevin Meng, David Bau, Alex Andonian, and Yonatan Belinkov. Locating and editing factual associations in gpt. *Advances in Neural Information Processing Systems*, 35:17359–17372, 2022.

- [24] Sewon Min, Mike Lewis, Luke Zettlemoyer, and Hannaneh Hajishirzi. Metaicl: Learning to learn in context. *arXiv preprint arXiv:2110.15943*, 2021.
- [25] OpenAI. Gpt-4 technical report, 2023.
- [26] Reese Pathak, Rajat Sen, Weihao Kong, and Abhimanyu Das. Transformers can optimally learn regression mixture models. *arXiv preprint arXiv:2311.08362*, 2023.
- [27] Kaare Brandt Petersen, Michael Syskind Pedersen, et al. The matrix cookbook. *Technical University of Denmark*, 7(15):510, 2008.
- [28] Alec Radford, Karthik Narasimhan, Tim Salimans, Ilya Sutskever, et al. Improving language understanding by generative pre-training. 2018.
- [29] Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, Ilya Sutskever, et al. Language models are unsupervised multitask learners. *OpenAI blog*, 1(8):9, 2019.
- [30] Allan Raventós, Mansheej Paul, Feng Chen, and Surya Ganguli. Pretraining task diversity and the emergence of non-bayesian in-context learning for regression. *arXiv preprint arXiv:2306.15063*, 2023.
- [31] George AF Seber. *A matrix handbook for statisticians*. John Wiley & Sons, 2008.
- [32] Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, et al. Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971*, 2023.
- [33] Alexander Tsigler and Peter L Bartlett. Benign overfitting in ridge regression. *J. Mach. Learn. Res.*, 24:123–1, 2023.
- [34] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. *Advances in neural information processing systems*, 30, 2017.
- [35] Johannes Von Oswald, Eyvind Niklasson, Ettore Randazzo, João Sacramento, Alexander Mordvintsev, Andrey Zhmoginov, and Max Vladymyrov. Transformers learn in-context by gradient descent. In *International Conference on Machine Learning*, pages 35151–35174. PMLR, 2023.
- [36] Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Rémi Louf, Morgan Funtowicz, Joe Davison, Sam Shleifer, Patrick von Platen, Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven Le Scao, Sylvain Gugger, Mariama Drame, Quentin Lhoest, and Alexander M. Rush. Huggingface’s transformers: State-of-the-art natural language processing, 2020.
- [37] Jingfeng Wu, Difan Zou, Zixiang Chen, Vladimir Braverman, Quanquan Gu, and Peter L Bartlett. How many pretraining tasks are needed for in-context learning of linear regression? In *The Twelfth International Conference on Learning Representations*, 2024.
- [38] Ruiqi Zhang, Spencer Frei, and Peter L Bartlett. Trained transformers learn linear models in-context. *arXiv preprint arXiv:2306.09927*, 2023.

A Notation and Variable Transformation

In order to simplify the proof, let's first describe a variable transformation scheme for our model and data coming from Assumption 3.1.

Definition A.1 (Variable Transformation). We fix M as the length of the contexts. Recall that $\mathbf{X}, \mathbf{x}$ are, respectively, the features in the context and the query input, while $\tilde{\beta}$ and β^* are the true task parameter for the inference prompt and its expectation, respectively. Following the Assumption 3.1, we have $\tilde{\beta} \sim \mathcal{N}(\beta^*, \Psi)$. From the definition for multivariate Gaussian distribution, we know there exists a random vector $\tilde{\theta}$ such that

$$\tilde{\beta} = \beta^* + \Psi^{\frac{1}{2}} \tilde{\theta}, \quad (\text{A.1})$$

where

$$\tilde{\theta} \sim \mathcal{N}(0, \mathbf{I}_d). \quad (\text{A.2})$$

Recall the noise vector is defined as $\varepsilon = \mathbf{y} - \mathbf{X}\beta$ and is generated from $\varepsilon \sim \mathcal{N}(0, \sigma^2 \cdot \mathbf{I}_M)$.

Rank deficient case Note that, even when Ψ is rank-deficient, the variable transformation in (A.1) and (A.2) still hold. This can be seen from the definition of the multivariate Gaussian distribution (Def 20.11 and the discussion below in [31]): A vector $\tilde{\beta} \in \mathbb{R}^d$ with mean β^* and covariance matrix Ψ has a multivariate normal distribution, if it has the same distribution as $\tilde{\mathbf{A}}\tilde{\theta} + \beta^*$ where $\tilde{\mathbf{A}}$ is any $d \times m$ matrix satisfying $\tilde{\mathbf{A}}\tilde{\mathbf{A}}^\top = \Psi$ and $\tilde{\theta} \sim \mathcal{N}(0_m, \mathbf{I}_m)$. Here, we can recover the variable transformation above if we take $m = d$ and $\tilde{\mathbf{A}} = \Psi^{\frac{1}{2}}$.

Notation Before we delve into the detailed proof, let's repeat the notation part with some additional notations. We use lowercase bold letters to note vectors and uppercase bold letters to denote matrices and tensors. For a vector $\mathbf{x}$ and a positive semi-definite (PSD) matrix $\mathbf{A}$, we denote $\|\mathbf{x}\|_{\mathbf{A}} := \sqrt{\mathbf{x}^\top \mathbf{A} \mathbf{x}}$. We denote $\langle \cdot, \cdot \rangle$ as the inner product, which is defined as $\langle \mathbf{x}, \mathbf{y} \rangle := \mathbf{x}^\top \mathbf{y}$ for vectors and $\langle \mathbf{A}, \mathbf{B} \rangle := \text{tr}(\mathbf{A}\mathbf{B}^\top)$ for matrices. For a matrix $\mathbf{A}$, we denote $\|\mathbf{A}\|_{op}, \|\mathbf{A}\|_F$ as the operator norm and the Frobenius norm, respectively. We denote $\mathbf{A}[i]$ as the i -th row of the matrix $\mathbf{A}$, $\mathbf{A}_{m,n}$ as the (m, n) -th entry, and $\mathbf{A}_{-1,-1}$ as the right-bottom entry. We denote $\mathbf{0}_n, \mathbf{0}_{m \times n}, \mathbf{I}_n$ as the zero vector, zero matrix and identity matrix, respectively.

For a positive semi-definite matrix $\mathbf{A}$, we denote $\mathbf{A}^{\frac{1}{2}}$ for the principle square root of $\mathbf{A}$, which is defined as the unique real matrix $\mathbf{B}$ such that $\mathbf{B}$ is positive semi-definite and $\mathbf{B}\mathbf{B} = \mathbf{B}\mathbf{B}^\top = \mathbf{A}$. For positive definite $\mathbf{A}$, its principle square root is also positive definite. We denote $\mathbf{A}^+$ as the Moore-Penrose pseudo-inverse for any matrix $\mathbf{A}$. We also denote $\mathbf{A}^{-\frac{1}{2}} = (\mathbf{A}^{\frac{1}{2}})^+$. We denote $\otimes$ as the Kronecker product. For compatible matrices of proper size $\mathbf{A}, \mathbf{B}, \mathbf{C}$, $\mathbf{B}^\top \otimes \mathbf{A}$ is a linear mapping which is defined by $(\mathbf{B}^\top \otimes \mathbf{A}) \circ \mathbf{C} = \mathbf{A}\mathbf{C}\mathbf{B}$. We denote $\Lambda := \Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}}$ and $\phi_1 \geq \phi_2 \geq \dots \geq \phi_d \geq 0$ are its ordered eigenvalues. We also denote $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_d \geq 0$ as the ordered eigenvalues of $\mathbf{H}$, and $\lambda_{-1} > 0$ as its minimal positive eigenvalue. Finally, we denote another important matrix

$$\Omega := \frac{M+1}{M} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} + \frac{\text{tr}(\mathbf{H}\Psi) + \sigma^2}{M} \cdot \mathbf{I}_d. \quad (\text{A.3})$$

Note that, under this definition and our assumption on $\mathbf{H}$ and Ψ , we have Ω is invertible.

B Proof of Theorem 4.1

Proof. The fact that $\inf_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f)$ does not depend on the vector β^* is subsumed in the Theorem 5.3, so we do not prove it here. We are going to prove the inequality in the theorem. First, from Theorem 5.3 we know that,

$$\inf_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f) = \sigma^2 + \text{tr}(\mathbf{H}\Psi) - \text{tr}\left(\left(\mathbf{H}^{\frac{1}{2}}\Psi\mathbf{H}^{\frac{1}{2}}\right)^2\Omega^{-1}\right), \quad (\text{B.1})$$

where Ω is defined in (A.3). Then, it suffices to lower bound $\inf_{f \in \mathcal{F}_{\text{LSA}}} \mathcal{R}(f)$. So let's take an arbitrary $f \in \mathcal{F}_{\text{LSA}}$, which is denoted as

$$f(\mathbf{E}) = \left[\mathbf{E} + \mathbf{W}_P^\top \mathbf{W}_V \mathbf{E} \mathbf{M} \frac{\mathbf{E}^\top \mathbf{W}_K^\top \mathbf{W}_Q \mathbf{E}}{M} \right]_{-1, -1},$$

where $\mathbf{E}$ is the token matrix defined in (3.1). Since the prediction is the right-bottom entry of the output matrix, we know only the last row of the product $\mathbf{W}_P^\top \mathbf{W}_V$ attends the prediction. Similarly, only the last column of the $\mathbf{E}$ on the far right in the above equation attends the prediction. Since the last column of $\mathbf{E}$ is $(\mathbf{x}^\top \ 0)^\top$, we know that only the first d rows of the product $\mathbf{W}_K^\top \mathbf{W}_Q$ enter the calculation (since other parts are multiplied by zero). Therefore, we denote

$$\mathbf{W}_P^\top \mathbf{W}_V = \begin{pmatrix} * & * \\ \mathbf{u}_{21}^\top & u_{-1} \end{pmatrix}, \quad \mathbf{W}_K^\top \mathbf{W}_Q = \begin{pmatrix} \mathbf{U}_{11} & * \\ \mathbf{u}_{12}^\top & * \end{pmatrix},$$

where $\mathbf{U}_{11} \in \mathbb{R}^{d \times d}$, $\mathbf{u}_{12}, \mathbf{u}_{21} \in \mathbb{R}^{d \times 1}$, $u_{-1} \in \mathbb{R}$, and $*$ denotes entries that do not enter the final prediction. The model prediction can be written as

$$\begin{aligned} f(\mathbf{E}) &= \begin{pmatrix} \mathbf{u}_{21}^\top & u_{-1} \end{pmatrix} \cdot \frac{\mathbf{E} \mathbf{M} \mathbf{E}^\top}{M} \cdot \begin{pmatrix} \mathbf{U}_{11} \\ \mathbf{u}_{12}^\top \end{pmatrix} \cdot \mathbf{x} \\ &= \left[\mathbf{u}_{21}^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{X} \cdot \mathbf{U}_{11} + \mathbf{u}_{21}^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{y} \cdot \mathbf{u}_{12}^\top + u_{-1} \cdot \frac{1}{M} \mathbf{y}^\top \mathbf{X} \cdot \mathbf{U}_{11} + u_{-1} \cdot \frac{1}{M} \mathbf{y}^\top \mathbf{y} \cdot \mathbf{u}_{12}^\top \right] \cdot \mathbf{x}. \end{aligned}$$

Step 1: simplify the risk function. We use $\tilde{\beta}$ to denote the task parameter. From the Assumption 3.1 and Definition A.1, we have

$$\mathbf{y} = \mathbf{X}\tilde{\beta} + \varepsilon, \quad y = \langle \tilde{\beta}, \mathbf{x} \rangle + \varepsilon, \quad \tilde{\beta} \sim \mathcal{N}(\beta^*, \Psi), \quad \tilde{\beta} = \beta^* + \Psi^{\frac{1}{2}} \tilde{\theta};$$

and

$$\mathbf{X}[i], \mathbf{x} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(\mathbf{0}, \mathbf{H}), \quad \varepsilon[i], \varepsilon \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \quad \tilde{\theta} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d).$$

Then the model output can be written as

$$\begin{aligned} f(\mathbf{E}) &= \left[\mathbf{u}_{21}^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{X} \cdot \mathbf{U}_{11} + \mathbf{u}_{21}^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{y} \cdot \mathbf{u}_{12}^\top + u_{-1} \cdot \frac{1}{M} \mathbf{y}^\top \mathbf{X} \cdot \mathbf{U}_{11} + u_{-1} \cdot \frac{1}{M} \mathbf{y}^\top \mathbf{y} \cdot \mathbf{u}_{12}^\top \right] \cdot \mathbf{x} \\ &= \left[\left(\mathbf{u}_{21} + u_{-1} \tilde{\beta} \right)^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{X} \cdot \left(\mathbf{U}_{11} + \tilde{\beta} \mathbf{u}_{12}^\top \right) \right] \cdot \mathbf{x} \\ &\quad + \left[\mathbf{u}_{21}^\top \cdot \frac{1}{M} \mathbf{X}^\top \varepsilon \cdot \mathbf{u}_{12}^\top + u_{-1} \cdot \frac{1}{M} \varepsilon^\top \mathbf{X} \cdot \mathbf{U}_{11} + u_{-1} \cdot \frac{2}{M} \varepsilon^\top \mathbf{X} \tilde{\beta} \mathbf{u}_{12}^\top + \frac{1}{M} \varepsilon^\top \varepsilon \cdot u_{-1} \mathbf{u}_{12}^\top \right] \cdot \mathbf{x}. \end{aligned}$$

To simplify the presentation, we denote

$$\begin{aligned} z_1^\top &= \left(\mathbf{u}_{21} + u_{-1} \tilde{\beta} \right)^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{X} \cdot \left(\mathbf{U}_{11} + \tilde{\beta} \mathbf{u}_{12}^\top \right), \\ z_2^\top &= \mathbf{u}_{21}^\top \cdot \frac{1}{M} \mathbf{X}^\top \varepsilon \cdot \mathbf{u}_{12}^\top + u_{-1} \cdot \frac{1}{M} \varepsilon^\top \mathbf{X} \cdot \mathbf{U}_{11} + u_{-1} \cdot \frac{2}{M} \varepsilon^\top \mathbf{X} \tilde{\beta} \mathbf{u}_{12}^\top \\ z_3^\top &= \frac{1}{M} \varepsilon^\top \varepsilon \cdot u_{-1} \mathbf{u}_{12}^\top. \end{aligned}$$

Since $\mathbf{x}, \mathbf{X}, \varepsilon, \tilde{\beta}$ are independent, we have

$$\begin{aligned}\mathcal{R}(f) &= \mathbb{E} \left(f(\mathbf{E}) - \langle \tilde{\beta}, \mathbf{x} \rangle - \varepsilon \right)^2 \\ &= \sigma^2 + \mathbb{E} \left(f(\mathbf{E}) - \langle \tilde{\beta}, \mathbf{x} \rangle \right)^2 \quad (\varepsilon \text{ is independent from other variables and zero-mean}) \\ &= \mathbb{E} \left[\left\langle \mathbf{z}_1 + \mathbf{z}_2 + \mathbf{z}_3 - \tilde{\beta}, \mathbf{x} \right\rangle^2 \right] + \sigma^2 \\ &= \left\langle \mathbf{H}, \mathbb{E} \left(\mathbf{z}_1 + \mathbf{z}_2 + \mathbf{z}_3 - \tilde{\beta} \right) \left(\mathbf{z}_1 + \mathbf{z}_2 + \mathbf{z}_3 - \tilde{\beta} \right)^\top \right\rangle + \sigma^2.\end{aligned}$$

Note that $\mathbf{z}_1$ does not contain ε , $\mathbf{z}_2$ is a linear form of ε , and $\mathbf{z}_3$ is a quadratic form of ε . Using $\varepsilon \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I}_d)$, we have $\mathbb{E}[(\mathbf{z}_1 - \tilde{\beta}) \cdot \mathbf{z}_2^\top] = \mathbf{0}$ and $\mathbb{E}[\mathbf{z}_2 \mathbf{z}_3^\top] = \mathbf{0}$. Therefore, we have

$$\begin{aligned}\mathcal{R}(f) - \sigma^2 &= \underbrace{\left\langle \mathbf{H}, \mathbb{E} \left(\mathbf{z}_1 - \tilde{\beta} \right) \left(\mathbf{z}_1 - \tilde{\beta} \right)^\top \right\rangle}_{S_1} + \underbrace{\left\langle \mathbf{H}, \mathbb{E} \mathbf{z}_2 \mathbf{z}_2^\top \right\rangle}_{S_2} + \underbrace{\left\langle \mathbf{H}, \mathbb{E} \mathbf{z}_3 \mathbf{z}_3^\top \right\rangle}_{S_3} \\ &\quad + 2 \underbrace{\left\langle \mathbf{H}, \mathbb{E} \left(\mathbf{z}_1 - \tilde{\beta} \right) \mathbf{z}_3^\top \right\rangle}_{S_4}.\end{aligned}\tag{B.2}$$

Step 2: compute S_1 . By Lemma H.4 and that $\mathbf{X}$ is independent of all other random variables, we have

$$\begin{aligned}&\langle \mathbf{H}, \mathbb{E} \mathbf{z}_1 \mathbf{z}_1^\top \rangle \\ &= \mathbb{E} \text{tr} \left[\left(\mathbf{U}_{11} + \tilde{\beta} \mathbf{u}_{12}^\top \right)^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{X} \cdot \left(\mathbf{u}_{21} + u_{-1} \tilde{\beta} \right) \left(\mathbf{u}_{21} + u_{-1} \tilde{\beta} \right)^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{X} \cdot \left(\mathbf{U}_{11} + \tilde{\beta} \mathbf{u}_{12}^\top \right) \mathbf{H} \right] \\ &= \frac{M+1}{M} \mathbb{E}_{\tilde{\beta}} \text{tr} \left[\left(\mathbf{U}_{11} + \tilde{\beta} \mathbf{u}_{12}^\top \right)^\top \cdot \mathbf{H} \cdot \left(\mathbf{u}_{21} + u_{-1} \tilde{\beta} \right) \left(\mathbf{u}_{21} + u_{-1} \tilde{\beta} \right)^\top \cdot \mathbf{H} \cdot \left(\mathbf{U}_{11} + \tilde{\beta} \mathbf{u}_{12}^\top \right) \mathbf{H} \right] \\ &\quad + \frac{1}{M} \mathbb{E}_{\tilde{\beta}} \text{tr} \left[\text{tr} \left(\left(\mathbf{u}_{21} + u_{-1} \tilde{\beta} \right) \left(\mathbf{u}_{21} + u_{-1} \tilde{\beta} \right)^\top \mathbf{H} \right) \left(\mathbf{U}_{11} + \tilde{\beta} \mathbf{u}_{12}^\top \right)^\top \mathbf{H} \left(\mathbf{U}_{11} + \tilde{\beta} \mathbf{u}_{12}^\top \right) \mathbf{H} \right].\end{aligned}$$

We denote

$$\mathbf{b} := \mathbf{u}_{21} + u_{-1} \beta^* \in \mathbb{R}^d, \quad \mathbf{A} := \mathbf{U}_{11} + \beta^* \mathbf{u}_{12}^\top \in \mathbb{R}^{d \times d},\tag{B.3}$$

then applying $\tilde{\beta} = \beta^* + \Psi^{\frac{1}{2}} \tilde{\theta}$, we get

$$\begin{aligned}&\langle \mathbf{H}, \mathbb{E} \mathbf{z}_1 \mathbf{z}_1^\top \rangle \\ &= \frac{M+1}{M} \mathbb{E}_{\tilde{\theta} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)} \text{tr} \left[\left(\mathbf{A} + \Psi^{\frac{1}{2}} \tilde{\theta} \mathbf{u}_{12}^\top \right)^\top \cdot \mathbf{H} \cdot \left(\mathbf{b} + u_{-1} \Psi^{\frac{1}{2}} \tilde{\theta} \right) \left(\mathbf{b} + u_{-1} \Psi^{\frac{1}{2}} \tilde{\theta} \right)^\top \cdot \mathbf{H} \cdot \left(\mathbf{A} + \Psi^{\frac{1}{2}} \tilde{\theta} \mathbf{u}_{12}^\top \right) \mathbf{H} \right] \\ &\quad + \frac{1}{M} \mathbb{E}_{\tilde{\theta} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)} \text{tr} \left[\left(\mathbf{b} + u_{-1} \Psi^{\frac{1}{2}} \tilde{\theta} \right) \left(\mathbf{b} + u_{-1} \Psi^{\frac{1}{2}} \tilde{\theta} \right)^\top \mathbf{H} \right] \text{tr} \left[\left(\mathbf{A} + \Psi^{\frac{1}{2}} \tilde{\theta} \mathbf{u}_{12}^\top \right)^\top \mathbf{H} \left(\mathbf{A} + \Psi^{\frac{1}{2}} \tilde{\theta} \mathbf{u}_{12}^\top \right) \mathbf{H} \right].\end{aligned}$$

Note that the first and third moments of $\tilde{\theta}$ are zero. So the above equation only involves the zeroth, second, and fourth moments of $\tilde{\theta}$. Therefore we have

$$\begin{aligned}\langle \mathbf{H}, \mathbb{E} \mathbf{z}_1 \mathbf{z}_1^\top \rangle &= \underbrace{\frac{M+1}{M} \text{tr} \left(\mathbf{A}^\top \mathbf{H} \mathbf{b} \mathbf{b}^\top \mathbf{H} \mathbf{A} \mathbf{H} \right) + \frac{1}{M} \text{tr} \left(\mathbf{b} \mathbf{b}^\top \mathbf{H} \right) \text{tr} \left(\mathbf{A}^\top \mathbf{H} \mathbf{A} \mathbf{H} \right)}_{\text{The leading term}} + T_2 + T_4,\end{aligned}\tag{B.4}$$

where T_2 and T_4 are the second order and the fourth order term, respectively. More concretely, the second order term is

$$T_2 = \frac{M+1}{M} \mathbb{E} \text{tr} \left\{ u_{-1} \mathbf{A}^\top \mathbf{H} \mathbf{b} \tilde{\theta}^\top \Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\theta} \mathbf{u}_{12}^\top \mathbf{H} + u_{-1} \mathbf{u}_{12} \tilde{\theta}^\top \Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\theta} \mathbf{b}^\top \mathbf{H} \mathbf{A} \mathbf{H} \right\}$$

$$\begin{aligned}
& + u_{-1} \mathbf{u}_{12} \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \mathbf{b} \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \mathbf{A} \mathbf{H} + u_{-1} \mathbf{A}^\top \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \mathbf{b}^\top \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \mathbf{u}_{12}^\top \mathbf{H} \\
& + \mathbf{u}_{12} \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \mathbf{b} \mathbf{b}^\top \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \mathbf{u}_{12}^\top \mathbf{H} + u_{-1}^2 \mathbf{A}^\top \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \mathbf{A} \mathbf{H} \Big\} \\
& + \frac{1}{M} \mathbb{E} \Big\{ u_{-1}^2 \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \cdot \text{tr} \left(\mathbf{A}^\top \mathbf{H} \mathbf{A} \mathbf{H} \right) + \mathbf{b}^\top \mathbf{H} \mathbf{b} \cdot \text{tr} \left(\mathbf{u}_{12} \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \mathbf{u}_{12}^\top \mathbf{H} \right) \\
& + 2u_{-1} \mathbf{b}^\top \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \cdot \text{tr} \left(\mathbf{u}_{12} \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \mathbf{A} \mathbf{H} \right) + 2u_{-1} \mathbf{b}^\top \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \cdot \text{tr} \left(\mathbf{A}^\top \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \mathbf{u}_{12}^\top \mathbf{H} \right) \Big\} \\
& = \frac{M+1}{M} \Big\{ 2u_{-1} \text{tr} \left(\Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}} \right) \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{A}^\top \mathbf{H} \mathbf{b} + 2u_{-1} \mathbf{b}^\top \mathbf{H} \Psi \mathbf{H} \mathbf{A} \mathbf{H} \mathbf{u}_{12} + \mathbf{b}^\top \mathbf{H} \Psi \mathbf{H} \mathbf{b} \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} \\
& + u_{-1}^2 \text{tr} \left(\mathbf{A}^\top \mathbf{H} \Psi \mathbf{H} \mathbf{A} \mathbf{H} \right) \Big\} + \frac{1}{M} \Big\{ u_{-1}^2 \text{tr} \left(\Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}} \right) \cdot \text{tr} \left(\mathbf{A}^\top \mathbf{H} \mathbf{A} \mathbf{H} \right) \\
& + \text{tr} \left(\Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}} \right) \cdot \mathbf{b}^\top \mathbf{H} \mathbf{b} \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} + 4u_{-1} \cdot \mathbf{b}^\top \mathbf{H} \Psi^{\frac{1}{2}} \mathbf{H} \mathbf{A} \mathbf{H} \mathbf{u}_{12} \Big\} \\
& = \frac{2(M+1)}{M} u_{-1} \mathbf{b}^\top \left[\text{tr}(\mathbf{H} \Psi) \mathbf{H} + \mathbf{H} \Psi \mathbf{H} \right] \mathbf{A} \mathbf{H} \mathbf{u}_{12} + \mathbf{b}^\top \left(\frac{M+1}{M} \mathbf{H} \Psi \mathbf{H} + \frac{1}{M} \text{tr}(\mathbf{H} \Psi) \mathbf{H} \right) \mathbf{b} \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} \\
& + u_{-1}^2 \text{tr} \left(\mathbf{A}^\top \left(\frac{M+1}{M} \mathbf{H} \Psi \mathbf{H} + \frac{1}{M} \text{tr}(\mathbf{H} \Psi) \mathbf{H} \right) \mathbf{A} \mathbf{H} \right) + \frac{4}{M} u_{-1} \cdot \mathbf{b}^\top \mathbf{H} \Psi \mathbf{H} \mathbf{A} \mathbf{H} \mathbf{u}_{12}. \quad (\text{B.5})
\end{aligned}$$

The fourth order term is

$$\begin{aligned}
T_4 & = \frac{M+1}{M} \mathbb{E} \text{tr} \Big\{ u_{-1}^2 \mathbf{u}_{12} \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \mathbf{u}_{12}^\top \mathbf{H} \Big\} \\
& \quad + \frac{1}{M} \mathbb{E} \text{tr} \Big\{ u_{-1}^2 \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \cdot \text{tr} \left(\mathbf{u}_{12} \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \mathbf{u}_{12}^\top \mathbf{H} \right) \Big\} \\
& = \frac{M+2}{M} u_{-1}^2 \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} \cdot \mathbb{E} \left(\tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \right)^2 \\
& = \frac{M+2}{M} u_{-1}^2 \cdot \left(2\text{tr}(\mathbf{H} \Psi \mathbf{H} \Psi) + \text{tr}(\mathbf{H} \Psi)^2 \right) \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12}. \quad (\text{B.6})
\end{aligned}$$

For the cross term in S_1 , we have

$$\begin{aligned}
& \left\langle \mathbf{H}, \mathbb{E} z_1 \tilde{\boldsymbol{\beta}}^\top \right\rangle \\
& = \mathbb{E} \Big\{ \left(\mathbf{u}_{21} + u_{-1} \tilde{\boldsymbol{\beta}} \right)^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{X} \cdot \left(\mathbf{U}_{11} + \tilde{\boldsymbol{\beta}} \mathbf{u}_{12}^\top \right) \mathbf{H} \tilde{\boldsymbol{\beta}} \Big\} \quad (\text{B.7}) \\
& = \mathbb{E} \Big\{ \left(\mathbf{u}_{21} + u_{-1} \tilde{\boldsymbol{\beta}} \right)^\top \mathbf{H} \left(\mathbf{U}_{11} + \tilde{\boldsymbol{\beta}} \mathbf{u}_{12}^\top \right) \mathbf{H} \tilde{\boldsymbol{\beta}} \Big\} \quad (\text{From the distribution of } \mathbf{X}) \\
& = \mathbb{E} \Big\{ \left(\mathbf{b} + u_{-1} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \right)^\top \mathbf{H} \left(\mathbf{A} + \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \mathbf{u}_{12}^\top \right) \mathbf{H} \left(\boldsymbol{\beta}^* + \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \right) \Big\} \quad (\text{By (B.3)}) \\
& = \mathbf{b}^\top \mathbf{H} \mathbf{A} \mathbf{H} \boldsymbol{\beta}^* + \mathbb{E} \Big\{ u_{-1} \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \mathbf{u}_{12}^\top \mathbf{H} \boldsymbol{\beta}^* + u_{-1} \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{H} \mathbf{A} \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} + \mathbf{b}^\top \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \mathbf{u}_{12}^\top \mathbf{H} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \Big\} \\
& = \mathbf{b}^\top \mathbf{H} \mathbf{A} \mathbf{H} \boldsymbol{\beta}^* + u_{-1} \text{tr}(\mathbf{H} \Psi) \cdot \mathbf{u}_{12}^\top \mathbf{H} \boldsymbol{\beta}^* + u_{-1} \text{tr}(\mathbf{H} \mathbf{A} \mathbf{H} \Psi) + \mathbf{u}_{12}^\top \mathbf{H} \Psi \mathbf{H} \mathbf{b}, \quad (\text{B.8})
\end{aligned}$$

where the last row comes from the fact that $\tilde{\boldsymbol{\theta}} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$. Moreover, we have

$$\left\langle \mathbf{H}, \mathbb{E} \tilde{\boldsymbol{\beta}} \tilde{\boldsymbol{\beta}}^\top \right\rangle = \boldsymbol{\beta}^{*\top} \mathbf{H} \boldsymbol{\beta}^* + \text{tr}(\mathbf{H} \Psi). \quad (\text{B.9})$$

Combining (B.4), (B.5), (B.6), (B.8), (B.9), we have

$$\begin{aligned}
S_1 & = \frac{M+1}{M} \text{tr} \left(\mathbf{A}^\top \mathbf{H} \mathbf{b} \mathbf{b}^\top \mathbf{H} \mathbf{A} \mathbf{H} \right) + \frac{1}{M} \text{tr} \left(\mathbf{b} \mathbf{b}^\top \mathbf{H} \right) \text{tr} \left(\mathbf{A}^\top \mathbf{H} \mathbf{A} \mathbf{H} \right) \\
& \quad + \frac{M+2}{M} u_{-1}^2 \cdot \left(2\text{tr}(\mathbf{H} \Psi \mathbf{H} \Psi) + \text{tr}(\mathbf{H} \Psi)^2 \right) \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12}
\end{aligned}$$

$$\begin{aligned}
& + \frac{2(M+1)}{M} u_{-1} \mathbf{b}^\top [\text{tr}(\mathbf{H}\Psi)\mathbf{H} + \mathbf{H}\Psi\mathbf{H}] \mathbf{A}\mathbf{H}\mathbf{u}_{12} \\
& + u_{-1}^2 \text{tr} \left(\mathbf{A}^\top \left(\frac{M+1}{M} \mathbf{H}\Psi\mathbf{H} + \frac{1}{M} \text{tr}(\mathbf{H}\Psi)\mathbf{H} \right) \mathbf{A}\mathbf{H} \right) + \frac{4}{M} u_{-1} \cdot \mathbf{b}^\top \mathbf{H}\Psi\mathbf{H}\mathbf{A}\mathbf{H}\mathbf{u}_{12} \\
& - 2 \left[\mathbf{b}^\top \mathbf{H}\mathbf{A}\mathbf{H}\beta^* + u_{-1} \text{tr}(\mathbf{H}\Psi) \cdot \mathbf{u}_{12}^\top \mathbf{H}\beta^* + u_{-1} \text{tr}(\mathbf{H}\mathbf{A}\mathbf{H}\Psi) \right. \\
& \left. + \mathbf{u}_{12}^\top \mathbf{H}\Psi\mathbf{H}\mathbf{b} \right] + \beta^{*\top} \mathbf{H}\beta^* + \text{tr}(\mathbf{H}\Psi) + \mathbf{b}^\top \left(\frac{M+1}{M} \mathbf{H}\Psi\mathbf{H} + \frac{1}{M} \text{tr}(\mathbf{H}\Psi)\mathbf{H} \right) \mathbf{b} \cdot \mathbf{u}_{12}^\top \mathbf{H}\mathbf{u}_{12}.
\end{aligned} \tag{B.10}$$

Step 3: other terms. Let us compute S_2, S_3 and S_4 . Using definitions, we rewrite $\mathbf{z}_2$ as

$$\begin{aligned}
\mathbf{z}_2^\top &= \mathbf{u}_{21}^\top \cdot \frac{1}{M} \mathbf{X}^\top \boldsymbol{\varepsilon} \cdot \mathbf{u}_{12}^\top + u_{-1} \cdot \frac{1}{M} \boldsymbol{\varepsilon}^\top \mathbf{X} \cdot \mathbf{U}_{11} + u_{-1} \cdot \frac{2}{M} \boldsymbol{\varepsilon}^\top \mathbf{X} \tilde{\boldsymbol{\beta}} \mathbf{u}_{12}^\top \\
&= \mathbf{b}^\top \cdot \frac{1}{M} \mathbf{X}^\top \boldsymbol{\varepsilon} \cdot \mathbf{u}_{12}^\top + u_{-1} \cdot \frac{1}{M} \boldsymbol{\varepsilon}^\top \mathbf{X} \cdot \mathbf{A} + u_{-1} \cdot \frac{2}{M} \boldsymbol{\varepsilon}^\top \mathbf{X} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \mathbf{u}_{12}^\top.
\end{aligned} \tag{B.11}$$

Since $\boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I}_M)$, $\tilde{\boldsymbol{\theta}} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$ and they are independent, all terms in S_2 vanish except if the term contains even orders of $\tilde{\boldsymbol{\theta}}$ or $\boldsymbol{\varepsilon}$. So we have

$$\begin{aligned}
S_2 &= \langle \mathbf{H}, \mathbb{E} \mathbf{z}_2 \mathbf{z}_2^\top \rangle = \frac{1}{M^2} \mathbb{E} \left\{ \mathbf{b}^\top \mathbf{X}^\top \boldsymbol{\varepsilon} \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} \cdot \boldsymbol{\varepsilon}^\top \mathbf{X} \cdot \mathbf{b} + u_{-1}^2 \boldsymbol{\varepsilon}^\top \mathbf{X} \mathbf{A} \mathbf{H} \mathbf{A}^\top \mathbf{X}^\top \boldsymbol{\varepsilon} \right. \\
&\quad \left. + 2u_{-1} \cdot \boldsymbol{\varepsilon}^\top \mathbf{X} \mathbf{A} \mathbf{H} \mathbf{u}_{12} \cdot \boldsymbol{\varepsilon}^\top \mathbf{X} \mathbf{b} + 4u_{-1}^2 \cdot \boldsymbol{\varepsilon}^\top \mathbf{X} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} \cdot \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \mathbf{X}^\top \boldsymbol{\varepsilon} \right\} \\
&= \frac{\sigma^2}{M} \left\{ \mathbf{b}^\top \mathbf{H} \mathbf{b} \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} + u_{-1}^2 \text{tr}(\mathbf{H} \mathbf{A} \mathbf{H} \mathbf{A}^\top) + 2u_{-1} \mathbf{u}_{12}^\top \mathbf{H} \mathbf{A}^\top \mathbf{H} \mathbf{b} + 4u_{-1}^2 \text{tr}(\mathbf{H}\Psi) \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} \right\}
\end{aligned} \tag{B.12}$$

For the other two terms, we have

$$S_3 = \frac{1}{M^2} u_{-1}^2 \mathbb{E} \{ \boldsymbol{\varepsilon}^\top \boldsymbol{\varepsilon} \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} \boldsymbol{\varepsilon}^\top \boldsymbol{\varepsilon} \} = \frac{\sigma^4 (M+2)}{M} u_{-1}^2 \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} \tag{B.13}$$

and

$$\begin{aligned}
S_4 &= 2 \langle \mathbf{H}, \mathbb{E} (\mathbf{z}_1 - \tilde{\boldsymbol{\beta}}) \mathbf{z}_3^\top \rangle \\
&= 2 \mathbb{E} \left\{ \left[(\mathbf{u}_{21} + u_{-1} \tilde{\boldsymbol{\beta}})^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{X} \cdot (\mathbf{U}_{11} + \tilde{\boldsymbol{\beta}} \mathbf{u}_{12}^\top) - \tilde{\boldsymbol{\beta}}^\top \right] \mathbf{H} \cdot \frac{1}{M} \boldsymbol{\varepsilon}^\top \boldsymbol{\varepsilon} \cdot u_{-1} \mathbf{u}_{12}^\top \right\} \\
&= 2\sigma^2 u_{-1} \mathbb{E} \left\{ \left[(\mathbf{b} + u_{-1} \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}})^\top \cdot \mathbf{H} \cdot (\mathbf{A} + \Psi^{\frac{1}{2}} \tilde{\boldsymbol{\theta}} \mathbf{u}_{12}^\top) - \tilde{\boldsymbol{\beta}}^{*\top} - \tilde{\boldsymbol{\theta}}^\top \Psi^{\frac{1}{2}} \right] \mathbf{H} \mathbf{u}_{12} \right\} \\
&= 2\sigma^2 u_{-1} \left[\mathbf{b}^\top \mathbf{H} \mathbf{A} - \tilde{\boldsymbol{\beta}}^{*\top} \right] \mathbf{H} \mathbf{u}_{12} + 2\sigma^2 u_{-1}^2 \text{tr}(\mathbf{H}\Psi) \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12}.
\end{aligned} \tag{B.14}$$

Step 4: combine all parts. Combining the four parts (B.10), (B.12), (B.13) and (B.14), we have

$$\mathcal{R}(f) - \sigma^2 = S_1 + S_2 + S_3 + S_4$$

We remark that in the (B.10), (B.12), (B.13) and (B.14), the parameters are $\mathbf{A}, \mathbf{b}, u_{-1}$ and $\mathbf{u}_{12}$, instead of the original parameter $\mathbf{U}_{11}, \mathbf{u}_{12}, \mathbf{u}_{21}, u_{-1}$. From the definitions of $\mathbf{A}$ and $\mathbf{b}$, we know there is a bijective map between $(\mathbf{U}_{11}, \mathbf{u}_{12}, \mathbf{u}_{21}, u_{-1})$ and $(\mathbf{A}, \mathbf{b}, \mathbf{u}_{12}, u_{-1})$, so these two parameterizations are equivalent when computing the minimal risk achieved by a model in a hypothesis class. From the expression of S_1, S_2, S_3, S_4 , we can rewrite the risk as

$$\mathcal{R}(f) - \sigma^2$$

$$\begin{aligned}
&= \mathbf{b}^\top \left(\frac{M+1}{M} \mathbf{H} \Psi \mathbf{H} + \frac{1}{M} \text{tr}(\mathbf{H} \Psi) \mathbf{H} \right) \mathbf{b} \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} + \frac{\sigma^2}{M} \cdot \mathbf{b}^\top \mathbf{H} \mathbf{b} \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} \\
&\quad + u_{-1}^2 \text{tr} \left(\mathbf{A}^\top \left(\frac{M+1}{M} \mathbf{H} \Psi \mathbf{H} + \frac{1}{M} \text{tr}(\mathbf{H} \Psi) \mathbf{H} \right) \mathbf{A} \mathbf{H} \right) + \frac{\sigma^2}{M} \cdot \text{tr}(\mathbf{H} \mathbf{A} \mathbf{H} \mathbf{A}^\top) \\
&\quad + 2u_{-1} \mathbf{b}^\top \left[\frac{1}{M} \text{tr}(\mathbf{H} \Psi) \mathbf{H} + \frac{M+1}{M} \mathbf{H} \Psi \mathbf{H} + \frac{\sigma^2}{M} \cdot \mathbf{H} \right] \mathbf{A} \mathbf{H} \mathbf{u}_{12} - 2u_{-1} \text{tr}(\mathbf{H} \mathbf{A} \mathbf{H} \Psi) - 2\mathbf{u}_{12}^\top \mathbf{H} \Psi \mathbf{H} \mathbf{b} \\
&\quad \underbrace{\hspace{15em}}_{\text{I}} \\
&\quad + \frac{1}{M} \left[\mathbf{b}^\top \mathbf{H} \mathbf{b} \cdot \text{tr}(\mathbf{A}^\top \mathbf{H} \mathbf{A} \mathbf{H}) + 4u_{-1} \cdot \mathbf{b}^\top \mathbf{H} \Psi \mathbf{H} \mathbf{A} \mathbf{H} \mathbf{u}_{12} + 4u_{-1}^2 \cdot \text{tr}(\mathbf{H} \Psi \mathbf{H} \Psi) \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} \right] \\
&\quad \underbrace{\hspace{15em}}_{\text{II}} \\
&\quad + \mathbf{b}^\top \left(\frac{M+1}{M} \mathbf{H} \mathbf{A} \mathbf{H} \mathbf{A}^\top \mathbf{H} \right) \mathbf{b} - 2\mathbf{b}^\top \mathbf{H} \mathbf{A} \mathbf{H} \beta^* + 2u_{-1} \text{tr}(\mathbf{H} \Psi) \cdot \mathbf{b}^\top \mathbf{H} \mathbf{A} \mathbf{H} \mathbf{u}_{12} + 2\sigma^2 \mathbf{b}^\top \mathbf{H} \mathbf{A} \mathbf{H} (u_{-1} \mathbf{u}_{12}) \\
&\quad \underbrace{\hspace{15em}}_{\text{III}} \\
&\quad + \beta^{*\top} \mathbf{H} \beta^* - 2(\text{tr}(\mathbf{H} \Psi) + \sigma^2) \cdot \beta^{*\top} \mathbf{H} (u_{-1} \mathbf{u}_{12}) + \text{tr}(\mathbf{H} \Psi) \\
&\quad + u_{-1}^2 \cdot \left[\left(2\text{tr}(\mathbf{H} \Psi \mathbf{H} \Psi) + \frac{M+2}{M} \text{tr}(\mathbf{H} \Psi)^2 \right) \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} + \left(2 + \frac{4}{M} \right) \sigma^2 \text{tr}(\mathbf{H} \Psi) + \frac{M+2}{M} \sigma^4 \right] \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} \\
&\quad \underbrace{\hspace{15em}}_{\text{IV}} \\
&= \text{I} + \text{II} + \text{III} + \text{IV}, \tag{B.15}
\end{aligned}$$

where I, II, III, IV are defined as above.

Step 5: lower bound the risk function. Let's first define a new matrix, which by definition is invertible:

$$\Omega := \frac{M+1}{M} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} + \frac{\text{tr}(\mathbf{H} \Psi) + \sigma^2}{M} \mathbf{I}_d. \tag{B.16}$$

So we can write I as

$$\begin{aligned}
\text{I} &= \mathbf{b}^\top \mathbf{H}^{\frac{1}{2}} \Omega \mathbf{H}^{\frac{1}{2}} \mathbf{b} \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12} + u_{-1}^2 \text{tr} \left(\mathbf{A}^\top \mathbf{H}^{\frac{1}{2}} \Omega \mathbf{H}^{\frac{1}{2}} \mathbf{A} \mathbf{H} \right) + 2u_{-1} \mathbf{b}^\top \mathbf{H}^{\frac{1}{2}} \Omega \mathbf{H}^{\frac{1}{2}} \mathbf{A} \mathbf{H} \mathbf{u}_{12} \\
&\quad - 2u_{-1} \text{tr}(\mathbf{H} \mathbf{A} \mathbf{H} \Psi) - 2\mathbf{u}_{12}^\top \mathbf{H} \Psi \mathbf{H} \mathbf{b} \\
&= \text{tr} \left[\left(\mathbf{H}^{\frac{1}{2}} \mathbf{u}_{12} \mathbf{b}^\top \mathbf{H}^{\frac{1}{2}} + u_{-1} \mathbf{H}^{\frac{1}{2}} \mathbf{A}^\top \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \right) \Omega \right. \\
&\quad \left. \cdot \left(\mathbf{H}^{\frac{1}{2}} \mathbf{u}_{12} \mathbf{b}^\top \mathbf{H}^{\frac{1}{2}} + u_{-1} \mathbf{H}^{\frac{1}{2}} \mathbf{A}^\top \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \right)^\top \right] - \text{tr} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right) \\
&\geq -\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right) \\
&= -\text{tr} \left(\left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right)^2 \Omega^{-1} \right), \tag{B.17}
\end{aligned}$$

where the last line comes from the fact that Ω and $\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}}$ commute.

Next, we claim $\text{II} \geq 0$. To see this, it suffices to notice that

$$\begin{aligned}
4u_{-1} \cdot \mathbf{b}^\top \mathbf{H} \Psi \mathbf{H} \mathbf{A} \mathbf{H} \mathbf{u}_{12} &\geq -4 \left\| \mathbf{b}^\top \mathbf{H}^{\frac{1}{2}} \right\|_F \cdot \left\| u_{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right\|_F \cdot \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{A} \mathbf{H}^{\frac{1}{2}} \right\|_F \cdot \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{u}_{12} \right\|_F \\
&\geq -\left\| \mathbf{b}^\top \mathbf{H}^{\frac{1}{2}} \right\|_F^2 \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{A} \mathbf{H}^{\frac{1}{2}} \right\|_F^2 - 4 \left\| u_{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right\|_F^2 \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{u}_{12} \right\|_F^2 \\
&\geq -\mathbf{b}^\top \mathbf{H} \mathbf{b} \cdot \text{tr}(\mathbf{A}^\top \mathbf{H} \mathbf{A} \mathbf{H}) - 4u_{-1}^2 \cdot \text{tr}(\mathbf{H} \Psi \mathbf{H} \Psi) \cdot \mathbf{u}_{12}^\top \mathbf{H} \mathbf{u}_{12}, \tag{B.18}
\end{aligned}$$

where the last line comes from the fact that $\|\mathbf{A}\|_F^2 = \text{tr}(\mathbf{A} \mathbf{A}^\top)$ for any matrix $\mathbf{A}$.

Then, let's consider III. Notice that actually, III can be viewed as a quadratic function of $\mathbf{A}^\top \mathbf{H} \mathbf{b}$, which can be easily minimized. More concretely, we have

$$\begin{aligned} & \text{III} + \frac{M}{M+1} \left(\beta^* - (\text{tr}(\mathbf{H}\Psi) + \sigma^2) u_{-1} \mathbf{u}_{12} \right)^\top \mathbf{H} \left(\beta^* - (\text{tr}(\mathbf{H}\Psi) + \sigma^2) u_{-1} \mathbf{u}_{12} \right) \\ &= \left[\mathbf{A}^\top \mathbf{H} \mathbf{b} - \frac{M}{M+1} \left(\beta^* - (\text{tr}(\mathbf{H}\Psi) + \sigma^2) u_{-1} \mathbf{u}_{12} \right) \right]^\top \left(\frac{M+1}{M} \mathbf{H} \right) \\ & \quad \cdot \left[\mathbf{A}^\top \mathbf{H} \mathbf{b} - \frac{M}{M+1} \left(\beta^* - (\text{tr}(\mathbf{H}\Psi) + \sigma^2) u_{-1} \mathbf{u}_{12} \right) \right] \\ &\geq 0. \end{aligned} \tag{B.19}$$

Therefore, one has

$$\text{III} \geq -\frac{M}{M+1} \left(\beta^* - (\text{tr}(\mathbf{H}\Psi) + \sigma^2) u_{-1} \mathbf{u}_{12} \right)^\top \mathbf{H} \left(\beta^* - (\text{tr}(\mathbf{H}\Psi) + \sigma^2) u_{-1} \mathbf{u}_{12} \right).$$

Combining the three parts above, one has

$$\begin{aligned} & \mathcal{R}(f) - \sigma^2 \\ &\geq \text{IV} - \text{tr} \left(\left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right)^2 \Omega^{-1} \right) \\ & \quad - \frac{M}{M+1} \left(\beta^* - (\text{tr}(\mathbf{H}\Psi) + \sigma^2) u_{-1} \mathbf{u}_{12} \right)^\top \mathbf{H} \left(\beta^* - (\text{tr}(\mathbf{H}\Psi) + \sigma^2) u_{-1} \mathbf{u}_{12} \right) \\ &= \underbrace{\text{tr}(\mathbf{H}\Psi) - \text{tr} \left(\left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right)^2 \Omega^{-1} \right)}_{\inf_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f) - \sigma^2} + \frac{1}{M+1} \beta^{*\top} \mathbf{H} \beta^* - \frac{2(\text{tr}(\mathbf{H}\Psi) + \sigma^2)}{M+1} \beta^{*\top} \mathbf{H} (u_{-1} \mathbf{u}_{12}) \\ & \quad + \left[2\text{tr}(\mathbf{H}\Psi \mathbf{H}\Psi) + \frac{3M+2}{M(M+1)} (\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2 \right] (u_{-1} \mathbf{u}_{12})^\top \mathbf{H} (u_{-1} \mathbf{u}_{12}). \end{aligned}$$

Therefore, one has

$$\begin{aligned} \mathcal{R}(f) - \inf_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f) &\geq \frac{1}{M+1} \beta^{*\top} \mathbf{H} \beta^* - \frac{2(\text{tr}(\mathbf{H}\Psi) + \sigma^2)}{M+1} \beta^{*\top} \mathbf{H} (u_{-1} \mathbf{u}_{12}) \\ & \quad + \left[2\text{tr}(\mathbf{H}\Psi \mathbf{H}\Psi) + \frac{3M+2}{M(M+1)} (\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2 \right] (u_{-1} \mathbf{u}_{12})^\top \mathbf{H} (u_{-1} \mathbf{u}_{12}). \end{aligned}$$

The right hand side in the above inequality is a quadratic function of $u_{-1} \mathbf{u}_{12}$, so we can take its global minimizer:

$$u_{-1} \mathbf{u}_{12} = \frac{\frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)}{M+1}}{2\text{tr}(\mathbf{H}\Psi \mathbf{H}\Psi) + \frac{3M+2}{M(M+1)} (\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2} \beta^*$$

to lower bound the risk gap as

$$\mathcal{R}(f) - \inf_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f) \geq \left[\frac{1}{M+1} - \frac{\frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2}{(M+1)^2}}{2\text{tr}(\mathbf{H}\Psi \mathbf{H}\Psi) + \frac{3M+2}{M(M+1)} (\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2} \right] \|\beta^*\|_{\mathbf{H}}^2.$$

Finally, to simplify the results, we notice that on the one hand,

$$\frac{\frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2}{(M+1)^2}}{2\text{tr}(\mathbf{H}\Psi \mathbf{H}\Psi) + \frac{3M+2}{M(M+1)} (\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2} \leq \frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2}{2(M+1)^2 \text{tr}(\mathbf{H}\Psi \mathbf{H}\Psi)}.$$

On the other hand, one has

$$\frac{\frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2}{(M+1)^2}}{2\text{tr}(\mathbf{H}\Psi\mathbf{H}\Psi) + \frac{3M+2}{M(M+1)}(\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2} \leq \frac{M}{(M+1)(3M+2)} \leq \frac{1}{3(M+1)}.$$

Therefore, we have

$$\mathcal{R}(f) - \inf_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f) \geq \max \left\{ \frac{2}{3(M+1)}, \frac{1}{M+1} - \frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2}{2(M+1)^2 \text{tr}(\mathbf{H}\Psi\mathbf{H}\Psi)} \right\} \cdot \|\beta^*\|_{\mathbf{H}}^2.$$

When $\frac{1}{M+1} - \frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2}{2(M+1)^2 \text{tr}(\mathbf{H}\Psi\mathbf{H}\Psi)} \geq \frac{2}{3(M+1)}$, we have $\frac{1}{M+1} \geq \frac{3(\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2}{2(M+1)^2 \text{tr}(\mathbf{H}\Psi\mathbf{H}\Psi)}$, which implies

$$\mathcal{R}(f) - \inf_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f) \geq \left[\frac{1}{M+1} - \frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2}{2(M+1)^2 \text{tr}(\mathbf{H}\Psi\mathbf{H}\Psi)} \right] \cdot \|\beta^*\|_{\mathbf{H}}^2 \geq \frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2}{(M+1)^2 \text{tr}(\mathbf{H}\Psi\mathbf{H}\Psi)} \cdot \|\beta^*\|_{\mathbf{H}}^2.$$

Therefore, we finally have

$$\mathcal{R}(f) - \inf_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f) \geq \max \left\{ \frac{2}{3(M+1)}, \frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2}{(M+1)^2 \text{tr}(\mathbf{H}\Psi\mathbf{H}\Psi)} \right\} \cdot \|\beta^*\|_{\mathbf{H}}^2.$$

Note that this holds for an arbitrary $f \in \mathcal{F}_{\text{LSA}}$, so the proof finishes by taking infimum on the left hand side. $\square$

C Proof of Theorem 5.2

Proof. Recall that from Assumption 3.1,

$$\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{H}), \quad y = \mathbf{x}^\top \tilde{\boldsymbol{\beta}} + \varepsilon, \quad \mathbf{X}[i] \sim \mathcal{N}(\mathbf{0}, \mathbf{H}), \quad \mathbf{y} = \mathbf{X} \tilde{\boldsymbol{\beta}} + \varepsilon,$$

where

$$\tilde{\boldsymbol{\beta}} \sim \mathcal{N}(\boldsymbol{\beta}^*, \boldsymbol{\Psi}), \quad \varepsilon \sim \mathcal{N}(0, \sigma^2), \quad \varepsilon \sim \mathcal{N}(\mathbf{0}, \sigma^2 \cdot \mathbf{I}_M).$$

Step 1: compute the risk function. From the independence of $\varepsilon, \boldsymbol{\beta}$ with other random variables, we have

$$\begin{aligned} \mathcal{R}_M(\boldsymbol{\beta}, \boldsymbol{\Gamma}) &= \mathbb{E} \left(\mathbf{x}^\top (\tilde{\boldsymbol{\beta}} - \boldsymbol{\beta}) + \varepsilon + \mathbf{x}^\top \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} (\boldsymbol{\beta} - \tilde{\boldsymbol{\beta}}) - \mathbf{x}^\top \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \varepsilon}{M} \right)^2 \\ &= \mathbb{E} \left(\mathbf{x}^\top \left(\mathbf{I}_d - \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} \right) \cdot (\tilde{\boldsymbol{\beta}} - \boldsymbol{\beta}) \right)^2 + \frac{1}{M^2} \mathbb{E} (\mathbf{x}^\top \boldsymbol{\Gamma} \mathbf{X}^\top \varepsilon)^2 + \sigma^2 \\ &= \left\langle \mathbf{H}, \mathbb{E} \left(\mathbf{I}_d - \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} \right)^{\otimes 2} \circ \mathbb{E} (\tilde{\boldsymbol{\beta}} - \boldsymbol{\beta})^{\otimes 2} \right\rangle + \frac{1}{M^2} \mathbb{E} (\mathbf{x}^\top \boldsymbol{\Gamma} \mathbf{X}^\top \varepsilon)^2 + \sigma^2, \end{aligned} \quad (\text{C.1})$$

where the last line comes from the independence between $\tilde{\boldsymbol{\beta}}$ and $\mathbf{X}, \mathbf{x}$, as well as the fact that $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{H})$ and the property of tensor product $\circ$. Note that, $\boldsymbol{\beta}$ and $\boldsymbol{\Gamma}$ are learnable parameters here, and we should set them apart from the task vector $\tilde{\boldsymbol{\beta}}$ or the prior mean vector $\tilde{\boldsymbol{\beta}}^*$. Recall that $\tilde{\boldsymbol{\beta}} \sim \mathcal{N}(\boldsymbol{\beta}^*, \boldsymbol{\Psi})$, we have $\mathbb{E} (\tilde{\boldsymbol{\beta}} - \boldsymbol{\beta})^{\otimes 2} = (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^{\otimes 2} + \boldsymbol{\Psi}$. Moreover, using the following

$$\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{H}), \quad \mathbf{X}[i] \sim \mathcal{N}(\mathbf{0}, \mathbf{H}), \quad \varepsilon \sim \mathcal{N}(\mathbf{0}, \sigma^2 \cdot \mathbf{I}_M),$$

we have that

$$\mathbb{E} (\mathbf{x}^\top \boldsymbol{\Gamma} \mathbf{X}^\top \varepsilon)^2 = \sigma^2 \mathbb{E} \text{tr} (\mathbf{X} \boldsymbol{\Gamma}^\top \mathbf{x} \mathbf{x}^\top \boldsymbol{\Gamma} \mathbf{X}^\top) = M \sigma^2 \cdot \langle \mathbf{H} \boldsymbol{\Gamma} \mathbf{H}, \boldsymbol{\Gamma}^\top \rangle.$$

Bridging the two terms above into (C.1), we get

$$\begin{aligned} \mathcal{R}_M(\boldsymbol{\beta}, \boldsymbol{\Gamma}) &= \left\langle \mathbf{H}, \mathbb{E} \left(\mathbf{I}_d - \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} \right)^{\otimes 2} \circ (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^{\otimes 2} \right\rangle \\ &\quad + \left\langle \mathbf{H}, \mathbb{E} \left(\mathbf{I}_d - \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} \right)^{\otimes 2} \circ \boldsymbol{\Psi} + \frac{\sigma^2}{M} \boldsymbol{\Gamma} \mathbf{H} \boldsymbol{\Gamma}^\top \right\rangle + \sigma^2 \\ &= \underbrace{\left\langle \mathbb{E} \left(\mathbf{I}_d - \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} \right)^{\otimes 2} \circ \mathbf{H}, (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^{\otimes 2} \right\rangle}_{V_1} \\ &\quad + \underbrace{\left\langle \mathbf{H}, \mathbb{E} \left(\left(\mathbf{I}_d - \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} \right)^\top \right)^{\otimes 2} \circ \boldsymbol{\Psi} + \frac{\sigma^2}{M} \boldsymbol{\Gamma} \mathbf{H} \boldsymbol{\Gamma}^\top \right\rangle}_{V_2} + \sigma^2 \end{aligned} \quad (\text{C.2})$$

First, let's compute V_1 . From the definition above, we have

$$V_1 = (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H}_\Gamma (\boldsymbol{\beta} - \boldsymbol{\beta}^*),$$

where

$$\mathbf{H}_\Gamma := \mathbb{E} \left(\left(\mathbf{I}_d - \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} \right)^\top \right)^{\otimes 2} \circ \mathbf{H} \quad (\text{C.3})$$

$$\begin{aligned}
&= \mathbb{E} \left(\mathbf{I}_d - \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} \right)^\top \mathbf{H} \left(\mathbf{I}_d - \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} \right) \quad (\text{the definition of the tensor product}) \\
&= \mathbf{H} - \mathbf{H} \left(\boldsymbol{\Gamma} + \boldsymbol{\Gamma}^\top \right) \mathbf{H} + \frac{\text{tr}(\mathbf{H} \boldsymbol{\Gamma}^\top \mathbf{H} \boldsymbol{\Gamma})}{M} \mathbf{H} + \frac{M+1}{M} \mathbf{H} \boldsymbol{\Gamma}^\top \mathbf{H} \boldsymbol{\Gamma} \mathbf{H} \\
&\quad (\mathbf{X}[i] \sim \mathcal{N}(\mathbf{0}, \mathbf{H}) \text{ and Lemma H.4}) \\
&= (\mathbf{I}_d - \boldsymbol{\Gamma} \mathbf{H})^\top \mathbf{H} (\mathbf{I}_d - \boldsymbol{\Gamma} \mathbf{H}) + \frac{\text{tr}(\mathbf{H} \boldsymbol{\Gamma}^\top \mathbf{H} \boldsymbol{\Gamma})}{M} \mathbf{H} + \frac{1}{M} \mathbf{H} \boldsymbol{\Gamma}^\top \mathbf{H} \boldsymbol{\Gamma} \mathbf{H} \succeq \mathbf{0}. \quad (\text{C.4})
\end{aligned}$$

Then, let's compute V_2 . We have

$$\begin{aligned}
&\mathbb{E} \left(\mathbf{I}_d - \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} \right)^{\otimes 2} \circ \boldsymbol{\Psi} + \frac{\sigma^2}{M} \boldsymbol{\Gamma} \mathbf{H} \boldsymbol{\Gamma}^\top = \mathbb{E} \left(\mathbf{I}_d - \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} \right) \boldsymbol{\Psi} \left(\mathbf{I}_d - \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} \right)^\top + \frac{\sigma^2}{M} \boldsymbol{\Gamma} \mathbf{H} \boldsymbol{\Gamma}^\top \\
&= \boldsymbol{\Psi} - \boldsymbol{\Gamma} \mathbf{H} \boldsymbol{\Psi} - \boldsymbol{\Psi} \mathbf{H} \boldsymbol{\Gamma}^\top + \boldsymbol{\Gamma} \left(\frac{\text{tr}(\mathbf{H} \boldsymbol{\Psi}) + \sigma^2}{M} \cdot \mathbf{H} + \frac{M+1}{M} \mathbf{H} \boldsymbol{\Psi} \mathbf{H} \right) \boldsymbol{\Gamma}^\top. \\
&\quad (\mathbf{X}[i] \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(\mathbf{0}, \mathbf{H}) \text{ and the Lemma H.4})
\end{aligned}$$

From the definition of $\boldsymbol{\Omega}$ in (A.3), we know that it is invertible. Moreover, it holds that

$$\begin{aligned}
&\mathbb{E} \left(\mathbf{I}_d - \boldsymbol{\Gamma} \cdot \frac{\mathbf{X}^\top \mathbf{X}}{M} \right)^{\otimes 2} \circ \boldsymbol{\Psi} + \frac{\sigma^2}{M} \boldsymbol{\Gamma} \mathbf{H} \boldsymbol{\Gamma}^\top = \boldsymbol{\Psi} - \boldsymbol{\Gamma} \mathbf{H} \boldsymbol{\Psi} - \boldsymbol{\Psi} \mathbf{H} \boldsymbol{\Gamma}^\top + \boldsymbol{\Gamma} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Gamma}^\top \\
&= \left(\boldsymbol{\Gamma} \mathbf{H}^{\frac{1}{2}} - \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \right) \boldsymbol{\Omega} \left(\boldsymbol{\Gamma} \mathbf{H}^{\frac{1}{2}} - \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \right)^\top + \boldsymbol{\Psi} - \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi}.
\end{aligned}$$

Therefore, we have

$$\begin{aligned}
V_2 &= \text{tr} \left[\left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Gamma} \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \right) \boldsymbol{\Omega} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Gamma} \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \right)^\top \right] \\
&\quad + \text{tr}(\mathbf{H} \boldsymbol{\Psi}) - \text{tr} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \right)
\end{aligned}$$

Step 2: solve the global minimum. Now we have

$$\begin{aligned}
&\mathcal{R}(\boldsymbol{\beta}, \boldsymbol{\Gamma}) - \sigma^2 \\
&= (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \left[(\mathbf{I}_d - \boldsymbol{\Gamma} \mathbf{H})^\top \mathbf{H} (\mathbf{I}_d - \boldsymbol{\Gamma} \mathbf{H}) + \frac{\text{tr}(\mathbf{H} \boldsymbol{\Gamma}^\top \mathbf{H} \boldsymbol{\Gamma})}{M} \mathbf{H} + \frac{1}{M} \mathbf{H} \boldsymbol{\Gamma}^\top \mathbf{H} \boldsymbol{\Gamma} \mathbf{H} \right] (\boldsymbol{\beta} - \boldsymbol{\beta}^*) \\
&\quad + \text{tr} \left[\left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Gamma} \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \right) \boldsymbol{\Omega} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Gamma} \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \right)^\top \right] \\
&\quad + \text{tr}(\mathbf{H} \boldsymbol{\Psi}) - \text{tr} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \right) \quad (\text{C.5}) \\
&\geq \text{tr}(\mathbf{H} \boldsymbol{\Psi}) - \text{tr} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \right) \\
&= \text{tr}(\mathbf{H} \boldsymbol{\Psi}) - \text{tr} \left(\left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \right)^2 \boldsymbol{\Omega}^{-1} \right),
\end{aligned}$$

where the last line comes from the fact that $\boldsymbol{\Omega}$ and $\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}}$ commute. Taking infimum over all $\boldsymbol{\beta}, \boldsymbol{\Gamma}$ in the left hand side, we have

$$\inf_{f \in \mathcal{F}_{\text{GD-}\boldsymbol{\beta}}} \mathcal{R}(f) - \sigma^2 = \inf_{\boldsymbol{\beta}, \boldsymbol{\Gamma}} \mathcal{R}(\boldsymbol{\beta}, \boldsymbol{\Gamma}) - \sigma^2 \geq \text{tr}(\mathbf{H} \boldsymbol{\Psi}) - \text{tr} \left(\left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \right)^2 \boldsymbol{\Omega}^{-1} \right).$$

On the other hand, we take

$$\beta = \beta^*, \quad \Gamma = \Gamma^* := \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \mathbf{H}^{-\frac{1}{2}}, \quad (\text{C.6})$$

where $\mathbf{H}^{\frac{1}{2}}$ is the principle square root of $\mathbf{H}$ and $\mathbf{H}^{-\frac{1}{2}}$ is its Moore-Penrose pseudo inverse (see notation part in Section A). Then, we have

$$\begin{aligned} & \mathcal{R}(\beta^*, \Gamma^*) - \sigma^2 \\ &= \text{tr} \left[\left(\mathbf{H}^{\frac{1}{2}} \Gamma^* \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \right) \Omega \left(\mathbf{H}^{\frac{1}{2}} \Gamma^* \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \right)^\top \right] \\ &+ \text{tr}(\mathbf{H} \Psi) - \text{tr} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right). \end{aligned}$$

Since

$$\begin{aligned} \mathbf{H}^{\frac{1}{2}} \Gamma^* \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} &= \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \\ &= \Omega^{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} - \Omega^{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}} \\ &\quad (\Omega \text{ and } \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \text{ commute}) \\ &= \Omega^{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} - \Omega^{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} = \mathbf{0}_{d \times d} \dots \\ &\quad (\mathbf{H}^{\frac{1}{2}} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}} = \mathbf{H}^{\frac{1}{2}}) \end{aligned}$$

Therefore, we have

$$\mathcal{R}(\beta^*, \Gamma^*) - \sigma^2 = \text{tr}(\mathbf{H} \Psi) - \text{tr} \left(\left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right)^2 \Omega^{-1} \right) \geq \inf_{f \in \mathcal{F}_{\text{GD}, \beta}} \mathcal{R}(f) - \sigma^2 = \inf_{\beta, \Gamma} \mathcal{R}(\beta, \Gamma) - \sigma^2.$$

Combining both directions, we conclude that

$$\inf_{f \in \mathcal{F}_{\text{GD}, \beta}} \mathcal{R}(f) - \sigma^2 = \inf_{\beta, \Gamma} \mathcal{R}(\beta, \Gamma) - \sigma^2 = \text{tr}(\mathbf{H} \Psi) - \text{tr} \left(\left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right)^2 \Omega^{-1} \right). \quad (\text{C.7})$$

Step 3: identify all global minimizers Above we show that $\mathcal{R}(\beta^*, \Gamma^*)$ achieves the global minimum of the ICL risk over GD- β class. Now we will figure out the sufficient and necessary condition to achieve the global minimal risk. From the proof above, we know that for any $\beta \in \mathbb{R}^d, \Gamma \in \mathbb{R}^{d \times d}$, it holds that

$$\mathcal{R}(\beta, \Gamma) = \inf_{f \in \mathcal{F}_{\text{GD}, \beta}} \mathcal{R}(f) \iff \begin{cases} V_1 = (\beta - \beta^*)^\top \mathbf{H}_\Gamma (\beta - \beta^*) = 0 \\ \text{tr} \left[\left(\mathbf{H}^{\frac{1}{2}} \Gamma \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \right) \Omega \left(\mathbf{H}^{\frac{1}{2}} \Gamma \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \right)^\top \right] = 0 \end{cases}$$

Since Ω is invertible, the second equation is equivalent to

$$\mathbf{H}^{\frac{1}{2}} \Gamma \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} = \mathbf{0}_{d \times d},$$

which is a linear system of Γ . Since Γ^* is proved to be one solution, all solutions of this equation is

$$\Gamma = \Gamma^* + \left\{ \mathbf{Z} \in \mathbb{R}^{d \times d} : \mathbf{H}^{\frac{1}{2}} \mathbf{Z} \mathbf{H}^{\frac{1}{2}} = \mathbf{0}_{d \times d} \right\} = \Gamma^* + \text{lm}(\mathbf{H}^{\otimes 2}),$$

where the last equation comes from Lemma H.6. Here, $+$ denotes Minkowski sum. This is defined as $A + B = \{a + b, a \in A, b \in B\}$ for two sets A, B and $a + B = \{a\} + B$ for an entry a and a set B . Under this condition, we know that

$$\mathbf{H}_\Gamma = \mathbf{H} - \mathbf{H} \left(\Gamma^* + \Gamma_M^{*\top} \right) \mathbf{H} + \frac{\text{tr}(\mathbf{H} \Gamma_M^{*\top} \mathbf{H} \Gamma^*)}{M} \mathbf{H} + \frac{M+1}{M} \mathbf{H} \Gamma_M^{*\top} \mathbf{H} \Gamma^* \mathbf{H}.$$

Consider the following inequality:

$$\mathbf{I}_d - \mathbf{H}^{\frac{1}{2}} \left(\Gamma^* + \Gamma^{*\top} \right) \mathbf{H}^{\frac{1}{2}} + \frac{\text{tr}(\mathbf{H} \Gamma_M^{*\top} \mathbf{H} \Gamma^*)}{M} \mathbf{I}_d + \frac{M+1}{M} \mathbf{H}^{\frac{1}{2}} \Gamma_M^{*\top} \mathbf{H} \Gamma^* \mathbf{H}^{\frac{1}{2}}$$

$$\begin{aligned}
& \preceq \mathbf{I}_d - \mathbf{H}^{\frac{1}{2}} \left(\mathbf{\Gamma}^* + \mathbf{\Gamma}^{*\top} \right) \mathbf{H}^{\frac{1}{2}} + \frac{M+1}{M} \mathbf{H}^{\frac{1}{2}} \mathbf{\Gamma}_M^{*\top} \mathbf{H} \mathbf{\Gamma}^* \mathbf{H}^{\frac{1}{2}} \\
& = \left(\sqrt{\frac{M}{M+1}} \mathbf{I}_d - \sqrt{\frac{M+1}{M}} \mathbf{H}^{\frac{1}{2}} \mathbf{\Gamma}^* \mathbf{H}^{\frac{1}{2}} \right) \cdot \left(\sqrt{\frac{M}{M+1}} \mathbf{I}_d - \sqrt{\frac{M+1}{M}} \mathbf{H}^{\frac{1}{2}} \mathbf{\Gamma}^* \mathbf{H}^{\frac{1}{2}} \right)^\top + \frac{M}{M+1} \mathbf{I}_d \\
& \succ \mathbf{0}_{d \times d}.
\end{aligned}$$

Therefore, we have

$$\begin{aligned}
& (\beta - \beta^*)^\top \mathbf{H}_\Gamma (\beta - \beta^*) = 0 \\
& \iff \left[(\beta - \beta^*)^\top \mathbf{H}^{\frac{1}{2}} \right] \left(\mathbf{I}_d - \mathbf{H}^{\frac{1}{2}} \left(\mathbf{\Gamma}^* + \mathbf{\Gamma}^{*\top} \right) \mathbf{H}^{\frac{1}{2}} + \frac{\text{tr} \left(\mathbf{H} \mathbf{\Gamma}_M^{*\top} \mathbf{H} \mathbf{\Gamma}^* \right)}{M} \mathbf{I}_d \right. \\
& \quad \left. + \frac{M+1}{M} \mathbf{H}^{\frac{1}{2}} \mathbf{\Gamma}_M^{*\top} \mathbf{H} \mathbf{\Gamma}^* \mathbf{H}^{\frac{1}{2}} \right) \left[\mathbf{H}^{\frac{1}{2}} (\beta - \beta^*) \right] = 0 \\
& \iff \mathbf{H}^{\frac{1}{2}} (\beta - \beta^*) = \mathbf{0}_d \iff \beta = \beta^* + \text{null} \left(\mathbf{H}^{\frac{1}{2}} \right) \iff \beta = \beta^* + \text{null} (\mathbf{H}),
\end{aligned}$$

where $\text{null}(\cdot)$ denotes the null space of a matrix and $+$ denotes the Minkowski sum. The last equivalence comes from Lemma H.6. Therefore, we conclude that the $f_{\beta, \Gamma}$ achieves the minimal ICL risk in GD- β class if and only if

$$\beta = \beta^* + \text{null} (\mathbf{H}), \quad \Gamma = \Gamma^* + \text{Im} (\mathbf{H}^{\otimes 2}).$$

Specially, if $\mathbf{H}$ is positive definite, there is unique global minimizer in GD- β class with parameters

$$\beta = \beta^*, \quad \Gamma = \Gamma^*.$$

Step 4: equivalence in the hypothesis class Finally, we show when $\mathbf{H}$ is rank-deficient, any global minimizer actually corresponds to one single function in $\mathcal{F}_{\text{GD-}\beta}$ almost surely. For an arbitrary global minimizer, we assume $\beta = \beta^* + \mathbf{h}$, $\Gamma = \Gamma^* + \mathbf{Z}$, where $\mathbf{h} \in \text{null} (\mathbf{H})$, $\mathbf{Z} \in \text{Im} (\mathbf{H}^{\otimes 2})$. Suppose we have a prompt $\mathbf{X}$, $\mathbf{y}$, $\mathbf{x}$, $\mathbf{y}$ which follows the Assumption 3.1 and the token matrix is formed by (3.1), we have

$$\begin{aligned}
f_{\beta, \Gamma} (\mathbf{E}) &= \left\langle \beta - \frac{\Gamma}{M} \mathbf{X}^\top (\mathbf{X} \beta - \mathbf{y}), \mathbf{x} \right\rangle \\
&= \left\langle \beta^* - \frac{\Gamma^*}{M} \mathbf{X}^\top (\mathbf{X} \beta^* - \mathbf{y}), \mathbf{x} \right\rangle \\
&\quad + \mathbf{h}^\top \mathbf{x} - \frac{1}{M} \mathbf{x}^\top \mathbf{Z} \mathbf{X}^\top \mathbf{X} \beta^* - \frac{1}{M} \mathbf{x}^\top \mathbf{Z} \mathbf{X}^\top \mathbf{X} \mathbf{h} - \frac{1}{M} \mathbf{x}^\top \mathbf{\Gamma}^* \mathbf{X}^\top \mathbf{X} \mathbf{h} + \mathbf{x}^\top \mathbf{Z} \mathbf{X}^\top \mathbf{y}.
\end{aligned}$$

Notice that $\mathbf{X}[i], \mathbf{x} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N} (\mathbf{0}_d, \mathbf{H})$, we have

$$\mathbb{E} (\mathbf{h}^\top \mathbf{x})^2 = \mathbf{h}^\top \mathbf{H} \mathbf{h} = 0$$

since $\mathbf{h} \in \text{null} (\mathbf{H})$. Therefore, we know $\mathbf{h}^\top \mathbf{x} = 0$ almost surely, and similarly, $\mathbf{X} \mathbf{h} = \mathbf{0}_d$ almost surely. Finally, we have

$$\mathbb{E} \left[\left(\frac{1}{M} \mathbf{x}^\top \mathbf{Z} \mathbf{X}^\top \right) \cdot \left(\frac{1}{M} \mathbf{x}^\top \mathbf{Z} \mathbf{X}^\top \right)^\top \right] = \frac{1}{M} \mathbb{E} \text{tr} [\mathbf{H} \mathbf{Z} \mathbf{H}] = 0,$$

which implies $\mathbf{x}^\top \mathbf{Z} \mathbf{X}^\top = \mathbf{0}_d$ almost surely. Therefore, we conclude

$$f_{\beta, \Gamma} (\mathbf{E}) = f_{\beta^*, \Gamma^*} (\mathbf{E}) \quad \text{almost surely.}$$

□

D Proof of Theorem 5.3

Proof. Recall the definition of Linear Transformer Block class:

$$f_{\text{LTB}} : \mathbb{R}^{(d+1) \times (M+1)} \rightarrow \mathbb{R}$$

$$\mathbf{E} \mapsto \left[\mathbf{W}_2^\top \mathbf{W}_1 \left(\mathbf{E} + \mathbf{W}_P^\top \mathbf{W}_V \mathbf{E} \mathbf{M} \frac{\mathbf{E}^\top \mathbf{W}_K^\top \mathbf{W}_Q \mathbf{E}}{M} \right) \right]_{-1, -1},$$

where $\mathbf{E}$ is the token matrix defined by (3.1). Let's take an arbitrary function f in the LTB class with trainable matrices

$$\mathbf{W}_K, \mathbf{W}_Q \in \mathbb{R}^{d_k \times (d+1)}, \quad \mathbf{W}_P, \mathbf{W}_V \in \mathbb{R}^{d_v \times (d+1)}, \quad \mathbf{W}_1, \mathbf{W}_2 \in \mathbb{R}^{d_f \times (d+1)}.$$

Similar to the proof of the Theorem 4.1, we know that only the last row of $\mathbf{W}_2^\top \mathbf{W}_1$ and the first d -columns of $\mathbf{W}_K^\top \mathbf{W}_Q$ attend the prediction. Therefore, we denote

$$\mathbf{W}_2^\top \mathbf{W}_1 = \begin{pmatrix} * & * \\ \gamma^\top & * \end{pmatrix}, \quad \mathbf{W}_2^\top \mathbf{W}_1 \mathbf{W}_P^\top \mathbf{W}_V = \begin{pmatrix} * & * \\ \mathbf{v}_{21}^\top & v_{-1} \end{pmatrix}, \quad \mathbf{W}_K^\top \mathbf{W}_Q = \begin{pmatrix} \mathbf{V}_{11} & * \\ \mathbf{v}_{12}^\top & * \end{pmatrix}, \quad (\text{D.1})$$

where $\gamma, \mathbf{v}_{12}, \mathbf{v}_{21}, \in \mathbb{R}^d, v_{-1} \in \mathbb{R}, \mathbf{V}_{11} \in \mathbb{R}^{d \times d}$, and $*$ denotes entries that do not enter the prediction. Then, the prediction of LTB function can be written as

$$f(\mathbf{E}) = \gamma^\top \mathbf{x} + \begin{pmatrix} \mathbf{v}_{21}^\top & v_{-1} \end{pmatrix} \cdot \frac{\mathbf{E} \mathbf{M} \mathbf{E}^\top}{M} \cdot \begin{pmatrix} \mathbf{V}_{11} \\ \mathbf{v}_{12}^\top \end{pmatrix} \cdot \mathbf{x}$$

$$= \left[\gamma^\top + \mathbf{v}_{21}^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{X} \cdot \mathbf{V}_{11} + \mathbf{v}_{21}^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{y} \cdot \mathbf{v}_{12}^\top + v_{-1} \cdot \frac{1}{M} \mathbf{y}^\top \mathbf{X} \cdot \mathbf{V}_{11} + v_{-1} \cdot \frac{1}{M} \mathbf{y}^\top \mathbf{y} \cdot \mathbf{v}_{12}^\top \right] \cdot \mathbf{x}.$$

Step 1: simplify the risk function. We use $\tilde{\beta}$ to denote the task parameter. From the Assumption 3.1 and Definition A.1, we have

$$\mathbf{y} = \mathbf{X} \tilde{\beta} + \varepsilon, \quad y = \langle \tilde{\beta}, \mathbf{x} \rangle + \varepsilon, \quad \tilde{\beta} \sim \mathcal{N}(\beta^*, \Psi), \quad \tilde{\beta} = \beta^* + \Psi^{\frac{1}{2}} \tilde{\theta};$$

and

$$\mathbf{X}[i], \mathbf{x} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(\mathbf{0}, \mathbf{H}), \quad \varepsilon[i], \varepsilon \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \quad \tilde{\theta} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d).$$

Then the model output can be written as

$$f(\mathbf{E}) = \left[\gamma^\top + \mathbf{v}_{21}^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{X} \cdot \mathbf{V}_{11} + \mathbf{v}_{21}^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{y} \cdot \mathbf{v}_{12}^\top + v_{-1} \cdot \frac{1}{M} \mathbf{y}^\top \mathbf{X} \cdot \mathbf{V}_{11} + v_{-1} \cdot \frac{1}{M} \mathbf{y}^\top \mathbf{y} \cdot \mathbf{v}_{12}^\top \right] \cdot \mathbf{x}$$

$$= \left[\gamma^\top + \left(\mathbf{v}_{21} + v_{-1} \tilde{\beta} \right)^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{X} \cdot \left(\mathbf{V}_{11} + \tilde{\beta} \mathbf{v}_{12}^\top \right) \right] \cdot \mathbf{x}$$

$$+ \left[\mathbf{v}_{21}^\top \cdot \frac{1}{M} \mathbf{X}^\top \varepsilon \cdot \mathbf{v}_{12}^\top + v_{-1} \cdot \frac{1}{M} \varepsilon^\top \mathbf{X} \cdot \mathbf{V}_{11} + v_{-1} \cdot \frac{2}{M} \varepsilon^\top \mathbf{X} \tilde{\beta} \mathbf{v}_{12}^\top + \frac{1}{M} \varepsilon^\top \varepsilon \cdot v_{-1} \mathbf{v}_{12}^\top \right] \cdot \mathbf{x}.$$

To simplify the presentation, we denote

$$\mathbf{z}_1^\top = \left(\mathbf{v}_{21} + v_{-1} \tilde{\beta} \right)^\top \cdot \frac{1}{M} \mathbf{X}^\top \mathbf{X} \cdot \left(\mathbf{V}_{11} + \tilde{\beta} \mathbf{v}_{12}^\top \right),$$

$$\mathbf{z}_2^\top = \mathbf{v}_{21}^\top \cdot \frac{1}{M} \mathbf{X}^\top \varepsilon \cdot \mathbf{v}_{12}^\top + v_{-1} \cdot \frac{1}{M} \varepsilon^\top \mathbf{X} \cdot \mathbf{V}_{11} + v_{-1} \cdot \frac{2}{M} \varepsilon^\top \mathbf{X} \tilde{\beta} \mathbf{v}_{12}^\top$$

$$\mathbf{z}_3^\top = \frac{1}{M} \varepsilon^\top \varepsilon \cdot v_{-1} \mathbf{v}_{12}^\top.$$

Since $\mathbf{x}, \mathbf{X}, \varepsilon, \tilde{\beta}$ are independent, we have

$$\mathcal{R}(f) = \mathbb{E} \left(f(\mathbf{E}) - \langle \tilde{\beta}, \mathbf{x} \rangle - \varepsilon \right)^2$$

$$\begin{aligned}
&= \sigma^2 + \mathbb{E} \left(f(\mathbf{E}) - \langle \tilde{\beta}, \mathbf{x} \rangle \right)^2 \quad (\varepsilon \text{ is independent from other variables and zero-mean}) \\
&= \mathbb{E} \left[\langle \mathbf{z}_1 + \mathbf{z}_2 + \mathbf{z}_3 + \gamma - \tilde{\beta}, \mathbf{x} \rangle^2 \right] + \sigma^2 \\
&= \left\langle \mathbf{H}, \mathbb{E} \left(\mathbf{z}_1 + \mathbf{z}_2 + \mathbf{z}_3 + \gamma - \tilde{\beta} \right) \left(\mathbf{z}_1 + \mathbf{z}_2 + \mathbf{z}_3 + \gamma - \tilde{\beta} \right)^\top \right\rangle + \sigma^2.
\end{aligned}$$

Note that $\mathbf{z}_1$ does not contain ε , $\mathbf{z}_2$ is a linear form of ε , and $\mathbf{z}_3$ is a quadratic form of ε . Using $\varepsilon \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I}_d)$, we have $\mathbb{E}[(\mathbf{z}_1 + \gamma - \tilde{\beta}) \cdot \mathbf{z}_2^\top] = \mathbf{0}$ and $\mathbb{E}[\mathbf{z}_2 \mathbf{z}_3^\top] = \mathbf{0}$. Therefore, we have

$$\begin{aligned}
\mathcal{R}(f) - \sigma^2 &= \underbrace{\left\langle \mathbf{H}, \mathbb{E} \left(\mathbf{z}_1 + \gamma - \tilde{\beta} \right) \left(\mathbf{z}_1 + \gamma - \tilde{\beta} \right)^\top \right\rangle}_{S'_1} + \underbrace{\left\langle \mathbf{H}, \mathbb{E} \mathbf{z}_2 \mathbf{z}_2^\top \right\rangle}_{S'_2} + \underbrace{\left\langle \mathbf{H}, \mathbb{E} \mathbf{z}_3 \mathbf{z}_3^\top \right\rangle}_{S'_3} \\
&\quad + 2 \underbrace{\left\langle \mathbf{H}, \mathbb{E} \left(\mathbf{z}_1 + \gamma - \tilde{\beta} \right) \mathbf{z}_3^\top \right\rangle}_{S'_4}. \tag{D.2}
\end{aligned}$$

Step 2: compute the risk function. Let's compute the risk function. Note that, the risk function is in the same form as the risk function of LSA layer, which is in the proof of Theorem 4.1 (see Section B and equation (B.2)). There are two differences: one is we replace $\mathbf{U}_{11}, \mathbf{u}_{12}, \mathbf{u}_{21}, u_{-1}$ here with $\mathbf{V}_{11}, \mathbf{v}_{12}, \mathbf{v}_{21}, v_{-1}$. The second difference is that we replace $\tilde{\beta}$ in (B.2) with $\tilde{\beta} - \gamma$. This is equivalent to replacing β^* with $\beta^* - \gamma$, since $\tilde{\beta} - \gamma \sim \mathcal{N}(\beta^* - \gamma, \Psi)$. Therefore, similar to (B.15), the risk function in (D.2) can be written as

$$\begin{aligned}
&\mathcal{R}(f) - \sigma^2 \\
&= \mathbf{b}^\top \left(\frac{M+1}{M} \mathbf{H} \Psi \mathbf{H} + \frac{1}{M} \text{tr}(\mathbf{H} \Psi) \mathbf{H} \right) \mathbf{b} \cdot \mathbf{v}_{12}^\top \mathbf{H} \mathbf{v}_{12} + \frac{\sigma^2}{M} \cdot \mathbf{b}^\top \mathbf{H} \mathbf{b} \cdot \mathbf{v}_{12}^\top \mathbf{H} \mathbf{v}_{12} \\
&\quad + v_{-1}^2 \text{tr} \left(\mathbf{A}^\top \left(\frac{M+1}{M} \mathbf{H} \Psi \mathbf{H} + \frac{1}{M} \text{tr}(\mathbf{H} \Psi) \mathbf{H} \right) \mathbf{A} \mathbf{H} \right) + \frac{\sigma^2}{M} \cdot \text{tr}(\mathbf{H} \mathbf{A} \mathbf{H} \mathbf{A}^\top) \\
&\quad + 2v_{-1} \mathbf{b}^\top \left[\frac{1}{M} \text{tr}(\mathbf{H} \Psi) \mathbf{H} + \frac{M+1}{M} \mathbf{H} \Psi \mathbf{H} + \frac{\sigma^2}{M} \cdot \mathbf{H} \right] \mathbf{A} \mathbf{H} \mathbf{v}_{12} - 2v_{-1} \text{tr}(\mathbf{H} \mathbf{A} \mathbf{H} \Psi) - 2\mathbf{v}_{12}^\top \mathbf{H} \Psi \mathbf{H} \mathbf{b} \\
&\quad \underbrace{\hspace{15em}}_{\text{V}} \\
&\quad + \frac{1}{M} \left[\mathbf{b}^\top \mathbf{H} \mathbf{b} \cdot \text{tr}(\mathbf{A}^\top \mathbf{H} \mathbf{A} \mathbf{H}) + 4v_{-1} \cdot \mathbf{b}^\top \mathbf{H} \Psi \mathbf{H} \mathbf{A} \mathbf{H} \mathbf{v}_{12} + 4v_{-1}^2 \cdot \text{tr}(\mathbf{H} \Psi \mathbf{H} \Psi) \cdot \mathbf{v}_{12}^\top \mathbf{H} \mathbf{v}_{12} \right] \\
&\quad \underbrace{\hspace{15em}}_{\text{VI}} \\
&\quad + \mathbf{b}^\top \left(\frac{M+1}{M} \mathbf{H} \mathbf{A} \mathbf{H} \mathbf{A}^\top \mathbf{H} \right) \mathbf{b} - 2\mathbf{b}^\top \mathbf{H} \mathbf{A} \mathbf{H} (\beta^* - \gamma) + 2v_{-1} \text{tr}(\mathbf{H} \Psi) \cdot \mathbf{b}^\top \mathbf{H} \mathbf{A} \mathbf{H} \mathbf{v}_{12} + 2\sigma^2 \mathbf{b}^\top \mathbf{H} \mathbf{A} \mathbf{H} (v_{-1} \mathbf{v}_{12}) \\
&\quad \underbrace{\hspace{15em}}_{\text{VII}} \\
&\quad + (\beta^* - \gamma)^\top \mathbf{H} (\beta^* - \gamma) - 2(\text{tr}(\mathbf{H} \Psi) + \sigma^2) \cdot (\beta^* - \gamma)^\top \mathbf{H} (v_{-1} \mathbf{v}_{12}) + \text{tr}(\mathbf{H} \Psi) \\
&\quad + v_{-1}^2 \cdot \left[\left(2\text{tr}(\mathbf{H} \Psi \mathbf{H} \Psi) + \frac{M+2}{M} \text{tr}(\mathbf{H} \Psi)^2 \right) \cdot \mathbf{v}_{12}^\top \mathbf{H} \mathbf{v}_{12} + \left(2 + \frac{4}{M} \right) \sigma^2 \text{tr}(\mathbf{H} \Psi) + \frac{M+2}{M} \sigma^4 \right] \cdot \mathbf{v}_{12}^\top \mathbf{H} \mathbf{v}_{12} \\
&\quad \underbrace{\hspace{15em}}_{\text{VIII}} \\
&= \text{V} + \text{VI} + \text{VII} + \text{VIII}, \tag{D.3}
\end{aligned}$$

where V, VI, VII, VIII are defined as above and $\Omega = \frac{M+1}{M} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} + \frac{\text{tr}(\mathbf{H} \Psi) + \sigma^2}{M} \cdot \mathbf{I}_d$, and $\mathbf{b} := \mathbf{v}_{21} + v_{-1} \beta^* \in \mathbb{R}^d$, $\mathbf{A} := \mathbf{V}_{11} + \beta^* \mathbf{v}_{12}^\top \in \mathbb{R}^{d \times d}$,

Step 3: solve the global minimum of the risk function. This is very similar to step 5 in Appendix B. The V term in (D.3) is actually equal to the I term in (B.15), so we have

$$\begin{aligned} \mathbf{V} &= \text{tr} \left[\left(\mathbf{H}^{\frac{1}{2}} \mathbf{v}_{12} \mathbf{b}^{\top} \mathbf{H}^{\frac{1}{2}} + v_{-1} \mathbf{H}^{\frac{1}{2}} \mathbf{A}^{\top} \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \mathbf{\Omega}^{-1} \right) \mathbf{\Omega} \left(\mathbf{H}^{\frac{1}{2}} \mathbf{v}_{12} \mathbf{b}^{\top} \mathbf{H}^{\frac{1}{2}} + v_{-1} \mathbf{H}^{\frac{1}{2}} \mathbf{A}^{\top} \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \mathbf{\Omega}^{-1} \right)^{\top} \right] \\ &\quad - \text{tr} \left(\mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \mathbf{\Omega}^{-1} \mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \right) \\ &\geq -\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \mathbf{\Omega}^{-1} \mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \right) \\ &= -\text{tr} \left(\left(\mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \right)^2 \mathbf{\Omega}^{-1} \right), \end{aligned}$$

where the last line comes from the fact that $\mathbf{\Omega}$ and $\mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}}$ commute. For the same reason, we know that the VI term above is equal to the II term in (B.15), which implies $\text{VI} \geq 0$, since

$$\begin{aligned} 4v_{-1} \cdot \mathbf{b}^{\top} \mathbf{H} \mathbf{\Psi} \mathbf{H} \mathbf{A} \mathbf{H} \mathbf{v}_{12} &\geq -4 \left\| \mathbf{b}^{\top} \mathbf{H}^{\frac{1}{2}} \right\|_F \cdot \left\| v_{-1} \mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \right\|_F \cdot \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{A} \mathbf{H}^{\frac{1}{2}} \right\|_F \cdot \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{v}_{12} \right\|_F \\ &\geq -\left\| \mathbf{b}^{\top} \mathbf{H}^{\frac{1}{2}} \right\|_F^2 \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{A} \mathbf{H}^{\frac{1}{2}} \right\|_F^2 - 4 \left\| v_{-1} \mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \right\|_F^2 \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{v}_{12} \right\|_F^2 \\ &= -\mathbf{b}^{\top} \mathbf{H} \mathbf{b} \cdot \text{tr} \left(\mathbf{A}^{\top} \mathbf{H} \mathbf{A} \mathbf{H} \right) - 4v_{-1}^2 \cdot \text{tr} \left(\mathbf{H} \mathbf{\Psi} \mathbf{H} \mathbf{\Psi} \right) \cdot \mathbf{v}_{12}^{\top} \mathbf{H} \mathbf{v}_{12}, \quad (\text{D.4}) \end{aligned}$$

where the last line comes from the fact that $\|\mathbf{A}\|_F^2 = \text{tr}(\mathbf{A} \mathbf{A}^{\top})$ for any matrix $\mathbf{A}$. The term VII above is equal to III term in (B.15), except that we replace β^* with $\beta^* - \gamma$. Therefore, we have

$$\begin{aligned} \text{VII} &+ \frac{M}{M+1} \left(\beta^* - \gamma - (\text{tr}(\mathbf{H} \mathbf{\Psi}) + \sigma^2) v_{-1} \mathbf{v}_{12} \right)^{\top} \mathbf{H} \left(\beta^* - \gamma - (\text{tr}(\mathbf{H} \mathbf{\Psi}) + \sigma^2) v_{-1} \mathbf{v}_{12} \right) \\ &= \left[\mathbf{A}^{\top} \mathbf{H} \mathbf{b} - \frac{M}{M+1} \left(\beta^* - \gamma - (\text{tr}(\mathbf{H} \mathbf{\Psi}) + \sigma^2) v_{-1} \mathbf{v}_{12} \right) \right]^{\top} \left(\frac{M+1}{M} \mathbf{H} \right) \\ &\quad \cdot \left[\mathbf{A}^{\top} \mathbf{H} \mathbf{b} - \frac{M}{M+1} \left(\beta^* - \gamma - (\text{tr}(\mathbf{H} \mathbf{\Psi}) + \sigma^2) v_{-1} \mathbf{v}_{12} \right) \right] \\ &\geq 0, \quad (\text{D.5}) \end{aligned}$$

which implies

$$\text{VII} \geq -\frac{M}{M+1} \left(\beta^* - \gamma - (\text{tr}(\mathbf{H} \mathbf{\Psi}) + \sigma^2) v_{-1} \mathbf{v}_{12} \right)^{\top} \mathbf{H} \left(\beta^* - \gamma - (\text{tr}(\mathbf{H} \mathbf{\Psi}) + \sigma^2) v_{-1} \mathbf{v}_{12} \right).$$

Combining the three parts above, one has

$$\begin{aligned} &\mathcal{R}(f) - \sigma^2 \\ &\geq \text{VIII} - \text{tr} \left(\left(\mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \right)^2 \mathbf{\Omega}^{-1} \right) \\ &\quad - \frac{M}{M+1} \left(\beta^* - \gamma - (\text{tr}(\mathbf{H} \mathbf{\Psi}) + \sigma^2) v_{-1} \mathbf{v}_{12} \right)^{\top} \mathbf{H} \left(\beta^* - \gamma - (\text{tr}(\mathbf{H} \mathbf{\Psi}) + \sigma^2) v_{-1} \mathbf{v}_{12} \right) \\ &= \underbrace{\text{tr}(\mathbf{H} \mathbf{\Psi}) - \text{tr} \left(\left(\mathbf{H}^{\frac{1}{2}} \mathbf{\Psi} \mathbf{H}^{\frac{1}{2}} \right)^2 \mathbf{\Omega}^{-1} \right)}_{\inf_{f \in \mathcal{F}_{\text{GD}, \beta}} \mathcal{R}(f) - \sigma^2} + \frac{1}{M+1} (\beta^* - \gamma)^{\top} \mathbf{H} (\beta^* - \gamma) \\ &\quad - \frac{2(\text{tr}(\mathbf{H} \mathbf{\Psi}) + \sigma^2)}{M+1} (\beta^* - \gamma)^{\top} \mathbf{H} (v_{-1} \mathbf{v}_{12}) \\ &\quad + \left[2\text{tr}(\mathbf{H} \mathbf{\Psi} \mathbf{H} \mathbf{\Psi}) + \frac{3M+2}{M(M+1)} (\text{tr}(\mathbf{H} \mathbf{\Psi}) + \sigma^2)^2 \right] (v_{-1} \mathbf{v}_{12})^{\top} \mathbf{H} (v_{-1} \mathbf{v}_{12}). \end{aligned}$$

Here, the global minimum of GD- β class is taken from Theorem 5.2, whose proof does not rely on the proof here. Therefore, one has

$$\begin{aligned} \mathcal{R}(f) - \inf_{f \in \mathcal{F}_{\text{GD},\beta}} \mathcal{R}(f) &\geq \frac{1}{M+1} (\beta^* - \gamma)^\top \mathbf{H} (\beta^* - \gamma) - \frac{2(\text{tr}(\mathbf{H}\Psi) + \sigma^2)}{M+1} (\beta^* - \gamma)^\top \mathbf{H} (v_{-1} \mathbf{v}_{12}) \\ &\quad + \left[2\text{tr}(\mathbf{H}\Psi\mathbf{H}\Psi) + \frac{3M+2}{M(M+1)} (\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2 \right] (v_{-1} \mathbf{v}_{12})^\top \mathbf{H} (v_{-1} \mathbf{v}_{12}). \end{aligned}$$

The right hand side in the above inequality is a quadratic function of $v_{-1} \mathbf{v}_{12}$, so we can take its global minimizer:

$$v_{-1} \mathbf{v}_{12} = \frac{\frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)}{M+1}}{2\text{tr}(\mathbf{H}\Psi\mathbf{H}\Psi) + \frac{3M+2}{M(M+1)} (\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2} (\beta^* - \gamma)$$

to lower bound the risk gap as

$$\mathcal{R}(f) - \inf_{f \in \mathcal{F}_{\text{GD},\beta}} \mathcal{R}(f) \geq \left[\frac{1}{M+1} - \frac{\frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2}{(M+1)^2}}{2\text{tr}(\mathbf{H}\Psi\mathbf{H}\Psi) + \frac{3M+2}{M(M+1)} (\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2} \right] \|\beta^* - \gamma\|_{\mathbf{H}}^2 \geq 0.$$

Note that, taking infimum on the left hand side, we get

$$\inf_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f) - \inf_{f \in \mathcal{F}_{\text{GD},\beta}} \mathcal{R}(f) \geq 0.$$

On the other hand, in the main text we have showed that $\mathcal{F}_{\text{GD},\beta} \subset \mathcal{F}_{\text{LTB}}$, which implies

$$\inf_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f) - \inf_{f \in \mathcal{F}_{\text{GD},\beta}} \mathcal{R}(f) \leq 0.$$

Therefore, we conclude

$$\inf_{f \in \mathcal{F}_{\text{LTB}}} \mathcal{R}(f) = \inf_{f \in \mathcal{F}_{\text{GD},\beta}} \mathcal{R}(f) = \text{tr}(\mathbf{H}\Psi) - \text{tr} \left(\left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right)^2 \Omega^{-1} \right),$$

where the last equation is from Theorem 5.2.

Step 4: sufficient and necessary conditions for global minimizers. Let's now verify the conditions for the global minimizers. For a function in LTB class, the sufficient and necessary condition for it to be a global minimizer is that inequalities in the above step all hold, which are

$$\mathbf{H}^{\frac{1}{2}} \mathbf{v}_{12} \mathbf{b}^\top \mathbf{H}^{\frac{1}{2}} + v_{-1} \mathbf{H}^{\frac{1}{2}} \mathbf{A}^\top \mathbf{H}^{\frac{1}{2}} = \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1}, \quad (\text{D.6})$$

$$\left\| \mathbf{b}^\top \mathbf{H}^{\frac{1}{2}} \right\|_F^2 \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{A} \mathbf{H}^{\frac{1}{2}} \right\|_F^2 = 4 \left\| v_{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right\|_F^2 \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{v}_{12} \right\|_F^2, \quad (\text{D.7})$$

$$v_{-1} \cdot \mathbf{b}^\top \mathbf{H} \Psi \mathbf{H} \mathbf{A} \mathbf{H} \mathbf{v}_{12} = \left\| \mathbf{b}^\top \mathbf{H}^{\frac{1}{2}} \right\|_F \cdot \left\| v_{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right\|_F \cdot \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{A} \mathbf{H}^{\frac{1}{2}} \right\|_F \cdot \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{v}_{12} \right\|_F, \quad (\text{D.8})$$

$$\mathbf{H}^{\frac{1}{2}} \left[\mathbf{A}^\top \mathbf{H} \mathbf{b} - \frac{M}{M+1} \left(\beta^* - \gamma - (\text{tr}(\mathbf{H}\Psi) + \sigma^2) v_{-1} \mathbf{v}_{12} \right) \right] = \mathbf{0}_{d \times d}, \quad (\text{D.9})$$

$$v_{-1} \mathbf{v}_{12} = \frac{\frac{(\text{tr}(\mathbf{H}\Psi) + \sigma^2)}{M+1}}{2\text{tr}(\mathbf{H}\Psi\mathbf{H}\Psi) + \frac{3M+2}{M(M+1)} (\text{tr}(\mathbf{H}\Psi) + \sigma^2)^2} (\beta^* - \gamma) \quad (\text{D.10})$$

$$\|\beta^* - \gamma\|_{\mathbf{H}} = 0. \quad (\text{D.11})$$

Let's first verify the necessary conditions of the system defined above. From (D.11) and Lemma H.5, we know $\gamma \in \beta^* + \text{null} \left(\mathbf{H}^{\frac{1}{2}} \right) = \beta^* + \text{null}(\mathbf{H})$. Then, we have $v_{-1} \mathbf{v}_{12} \in \text{null}(\mathbf{H})$. Then, in (D.7), we know

$$\left\| \mathbf{b}^\top \mathbf{H}^{\frac{1}{2}} \right\|_F^2 \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{A} \mathbf{H}^{\frac{1}{2}} \right\|_F^2 = 4 \left\| v_{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right\|_F^2 \left\| \mathbf{H}^{\frac{1}{2}} \mathbf{v}_{12} \right\|_F^2 = 4 \left\| \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right\|_F^2 \left\| v_{-1} \mathbf{H}^{\frac{1}{2}} \mathbf{v}_{12} \right\|_F^2 = 0,$$

which implies either $\mathbf{H}^{\frac{1}{2}}\mathbf{b} = \mathbf{0}_d$ or $\mathbf{H}^{\frac{1}{2}}\mathbf{A}\mathbf{H}^{\frac{1}{2}} = \mathbf{0}_{d \times d}$. If we assume $\mathbf{H}^{\frac{1}{2}}\mathbf{A}\mathbf{H}^{\frac{1}{2}} = \mathbf{0}_{d \times d}$, then (D.6) implies $\mathbf{H}^{\frac{1}{2}}\mathbf{v}_{12}\mathbf{b}^\top \mathbf{H}^{\frac{1}{2}} = \mathbf{H}^{\frac{1}{2}}\Psi\mathbf{H}^{\frac{1}{2}}\Omega^{-1}$. From our assumption, we know $\text{rank}(\mathbf{H}^{\frac{1}{2}}\Psi\mathbf{H}^{\frac{1}{2}}) \geq 2$, which implies $\text{rank}(\mathbf{H}^{\frac{1}{2}}\Psi\mathbf{H}^{\frac{1}{2}}) \geq 2$, since for any matrix $\mathbf{Z}$, it holds that $\text{rank}(\mathbf{Z}) = \text{rank}(\mathbf{Z}\mathbf{Z}^\top)$. Since multiplication by an invertible matrix does not change the rank, we know $\text{rank}(\mathbf{H}^{\frac{1}{2}}\Psi\mathbf{H}^{\frac{1}{2}}\Omega^{-1}) \geq 2$, while $\mathbf{H}^{\frac{1}{2}}\mathbf{v}_{12}\mathbf{b}^\top \mathbf{H}^{\frac{1}{2}}$ is a matrix of rank at most one, which contradicts with (D.6). Therefore, we have

$$\mathbf{H}^{\frac{1}{2}}\mathbf{b} = \mathbf{0}, \quad (\text{D.12})$$

$$v_{-1}\mathbf{H}^{\frac{1}{2}}\mathbf{A}^\top \mathbf{H}^{\frac{1}{2}} = \mathbf{H}^{\frac{1}{2}}\Psi\mathbf{H}^{\frac{1}{2}}\Omega^{-1}. \quad (\text{D.13})$$

Recall in Theorem 5.2, we have defined

$$\Gamma^* := \Psi\mathbf{H}^{\frac{1}{2}}\Omega^{-1}\mathbf{H}^{-\frac{1}{2}}.$$

Simple calculation shows

$$\mathbf{H}^{\frac{1}{2}}\Gamma^*\mathbf{H}^{\frac{1}{2}} = \mathbf{H}^{\frac{1}{2}}\Psi\mathbf{H}^{\frac{1}{2}}\Omega^{-1}\mathbf{H}^{-\frac{1}{2}}\mathbf{H}^{\frac{1}{2}} = \Omega^{-1}\mathbf{H}^{\frac{1}{2}}\Psi\mathbf{H}^{\frac{1}{2}}\mathbf{H}^{-\frac{1}{2}}\mathbf{H}^{\frac{1}{2}} = \Omega^{-1}\mathbf{H}^{\frac{1}{2}}\Psi\mathbf{H}^{\frac{1}{2}} = \mathbf{H}^{\frac{1}{2}}\Psi\mathbf{H}^{\frac{1}{2}}\Omega^{-1},$$

where the second and the last equalities come from the fact that Ω and $\mathbf{H}^{\frac{1}{2}}\Psi\mathbf{H}^{\frac{1}{2}}$ commute, and the third equality comes from the property of Moor Penrose pseudo-inverse. Therefore, one solution of (D.13) is $v_{-1}\mathbf{A} = \Gamma^{*\top}$. Since (D.13) is a linear equation to $v_{-1}\mathbf{A}$, we know its full solution is

$$v_{-1}\mathbf{A} \in \Gamma^{*\top} + \left\{ \mathbf{Z} : \mathbf{H}^{\frac{1}{2}}\mathbf{Z}\mathbf{H}^{\frac{1}{2}} = \mathbf{0}_{d \times d} \right\} = \Gamma^{*\top} + \text{Im}(\mathbf{H}^{\otimes 2}).$$

The solution to (D.12) is

$$\mathbf{v}_{21} = -v_{-1}\beta^* + \text{null}(\mathbf{H}^{\frac{1}{2}}) = -v_{-1}\beta^* + \text{null}(\mathbf{H}).$$

Therefore, we know the necessary conditions for (D.6) to (D.11) are

$$\begin{cases} v_{-1} \neq 0, \\ v_{-1}\mathbf{v}_{12} \in \text{null}(\mathbf{H}), \\ \mathbf{v}_{21} = -v_{-1}\beta^* + \text{null}(\mathbf{H}), \\ v_{-1}\mathbf{V}_{11} \in \Gamma^{*\top} - v_{-1}\beta^*\mathbf{v}_{12}^\top + \text{Im}(\mathbf{H}^{\otimes 2}), \\ \gamma = \beta^* + \text{null}(\mathbf{H}) \end{cases} \quad (\text{D.14})$$

It is easy to verify these equations above are also sufficient conditions of (D.6) to (D.11) by directly replacing each variable with its value and validating equations from (D.6) to (D.11).

Specially, if $\mathbf{H}$ is positive definite and hence, invertible, the global minimizer is unique up to a scaling to v_{-1} :

$$v_{-1} \neq 0, \quad \mathbf{v}_{12} = \mathbf{0}_d, \quad \mathbf{v}_{21} = -v_{-1}\beta^*, \quad \mathbf{V}_{11} = \frac{1}{v_{-1}} \cdot \Gamma^{*\top}, \quad \gamma = \beta^*. \quad (\text{D.15})$$

Step 5: equivalence in the hypothesis class. Finally, we will prove that any global minimizer of $\mathcal{F}_{\text{LTB}}$ is actually equivalent to one single function in $\mathcal{F}_{\text{LTB}}$ almost surely. More concretely, let's take a function $f \in \mathcal{F}_{\text{LTB}}$ with parameters $\mathbf{W}_K, \mathbf{W}_Q, \mathbf{W}_P, \mathbf{W}_V, \mathbf{W}_1, \mathbf{W}_2$ and recall the parameter transformation in (D.1). We assume equations in (D.14) and Assumption 3.1 hold. Then, for any vector $\mathbf{a} \in \text{null}(\mathbf{H})$, one has $\mathbf{x}^\top \mathbf{a} = \mathbf{x}_i^\top \mathbf{a} = 0$ since $\mathbf{x}, \mathbf{x}_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(\mathbf{0}_d, \mathbf{H})$. We denote

$$\mathbf{v}_{21} = -v_{-1}\beta^* + \mathbf{a}_1, \quad v_{-1}\mathbf{V}_{11} = \Gamma^{*\top} - v_{-1}\beta^*\mathbf{v}_{12}^\top + \mathbf{Z}, \quad \gamma = \beta^* + \mathbf{a}_2,$$

where $\mathbf{a}_1, \mathbf{a}_2 \in \mathbf{H}, \mathbf{H}^{\frac{1}{2}}\mathbf{Z}\mathbf{H}^{\frac{1}{2}} = \mathbf{0}_{d \times d}$. Then, we have

$$f(\mathbf{E}) = \left[\mathbf{W}_2^\top \mathbf{W}_1 \left(\mathbf{E} + \mathbf{W}_P^\top \mathbf{W}_V \mathbf{E} \mathbf{M} \frac{\mathbf{E}^\top \mathbf{W}_K^\top \mathbf{W}_Q \mathbf{E}}{M} \right) \right]_{-1, -1}$$

$$\begin{aligned}
&= (\boldsymbol{\beta}^* + \mathbf{a}_2)^\top \mathbf{x} + \begin{pmatrix} -v_{-1}\boldsymbol{\beta}^{*\top} + \mathbf{a}_1^\top & v_{-1} \end{pmatrix} \cdot \frac{1}{M} \begin{pmatrix} \mathbf{X}^\top \mathbf{X} & \mathbf{X}^\top \mathbf{y} \\ \mathbf{y}^\top \mathbf{X} & \mathbf{y}^\top \mathbf{y} \end{pmatrix} \cdot \begin{pmatrix} \mathbf{V}_{11} \\ \mathbf{v}_{12}^\top \end{pmatrix} \cdot \mathbf{x} \\
&= \boldsymbol{\beta}^{*\top} \mathbf{x} + \begin{pmatrix} -\boldsymbol{\beta}^{*\top} & 1 \end{pmatrix} \cdot \frac{1}{M} \begin{pmatrix} \mathbf{X}^\top \mathbf{X} & \mathbf{X}^\top \mathbf{y} \\ \mathbf{y}^\top \mathbf{X} & \mathbf{y}^\top \mathbf{y} \end{pmatrix} \cdot \begin{pmatrix} v_{-1} \mathbf{V}_{11} \mathbf{x} \\ v_{-1} \mathbf{v}_{12}^\top \mathbf{x} \end{pmatrix} \quad (\mathbf{X} \mathbf{a}_1 = \mathbf{0}, \mathbf{a}_2^\top \mathbf{x} = 0) \\
&= \boldsymbol{\beta}^{*\top} \mathbf{x} + \begin{pmatrix} -\boldsymbol{\beta}^{*\top} & 1 \end{pmatrix} \cdot \frac{1}{M} \begin{pmatrix} \mathbf{X}^\top \mathbf{X} & \mathbf{X}^\top \mathbf{y} \\ \mathbf{y}^\top \mathbf{X} & \mathbf{y}^\top \mathbf{y} \end{pmatrix} \cdot \begin{pmatrix} \boldsymbol{\Gamma}^{*\top} \mathbf{x} \\ \mathbf{0}_d^\top \end{pmatrix} \quad (\mathbf{v}_{12}^\top \mathbf{x} = 0, \mathbf{X} \mathbf{V}_{11} \mathbf{x} = \mathbf{0}) \\
&= \left\langle \boldsymbol{\beta}^* - \frac{\boldsymbol{\Gamma}^*}{M} \mathbf{X}^\top (\mathbf{X} \boldsymbol{\beta} - \mathbf{y}), \mathbf{x} \right\rangle = f_{\boldsymbol{\beta}^*, \boldsymbol{\Gamma}^*}(\mathbf{E}),
\end{aligned}$$

where $f_{\boldsymbol{\beta}^*, \boldsymbol{\Gamma}^*}(\cdot)$ is the GD- $\boldsymbol{\beta}$ function defined in (5.1). □

E Proof of Corollary 6.2

Proof. We denote $\phi_1 \geq \phi_2 \geq \dots \geq \phi_d \geq 0$ are ordered eigenvalues of $\Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}}$. From Theorem 5.2, we have

$$\inf_{f \in \mathcal{F}_{\text{GD}, \beta}} \mathcal{R}(f) - \sigma^2 = \text{tr} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right) - \text{tr} \left(\left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right)^2 \Omega^{-1} \right),$$

where

$$\Omega := \frac{M+1}{M} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} + \frac{\text{tr}(\mathbf{H} \Psi) + \sigma^2}{M} \cdot \mathbf{I}_d.$$

Therefore, we have

$$\begin{aligned} \inf_{f \in \mathcal{F}_{\text{GD}, \beta}} \mathcal{R}(f) - \sigma^2 &= \text{tr} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \cdot \left(\Omega - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right) \Omega^{-1} \right) \\ &= \frac{1}{M} \text{tr} \left(\Omega^{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \cdot \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} + \left(\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right) + \sigma^2 \right) \cdot \mathbf{I}_d \right) \right). \end{aligned}$$

Since

$$\left(\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right) + \sigma^2 \right) \cdot \mathbf{I}_d \preceq \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} + \left(\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right) + \sigma^2 \right) \cdot \mathbf{I}_d \preceq 2 \left(\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right) + \sigma^2 \right) \cdot \mathbf{I}_d,$$

we have

$$\begin{aligned} \inf_{f \in \mathcal{F}_{\text{GD}, \beta}} \mathcal{R}(f) - \sigma^2 &\simeq \frac{\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right) + \sigma^2}{M} \text{tr} \left(\Omega^{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right) \\ &= \bar{\phi} \cdot \sum_{i=1}^d \frac{\phi_i}{\frac{M+1}{M} \phi_i + \bar{\phi}} \\ &\simeq \sum_{i=1}^d \min \{ \phi_i, \bar{\phi} \}. \end{aligned}$$

Therefore, we conclude. □

F Proof of Theorem 6.1

F.1 Proof of first equation

Proof. From the classical bias-variance decomposition, we know

$$\begin{aligned}\mathcal{L}(g; \mathbf{X}) &:= \mathbb{E} \left[(g(\mathbf{X}, \mathbf{y}, \mathbf{x}) - y)^2 \mid \mathbf{X} \right] \\ &= \mathbb{E} \left[(g(\mathbf{X}, \mathbf{y}, \mathbf{x}) - \mathbb{E}[y \mid \mathbf{X}, \mathbf{y}, \mathbf{x}])^2 \mid \mathbf{X} \right] + \mathbb{E} \left[(\mathbb{E}[y \mid \mathbf{X}, \mathbf{y}, \mathbf{x}] - y)^2 \mid \mathbf{X} \right]\end{aligned}$$

since the cross term vanishes. The second term does not depend on g and hence, the Bayesian optimal estimator is given by the posterior mean, i.e.,

$$\hat{y}_{\text{Bayes}} = \mathbb{E}[y \mid \mathbf{X}, \mathbf{y}, \mathbf{x}] = \left\langle \mathbb{E}[\tilde{\beta} \mid \mathbf{X}, \mathbf{y}, \mathbf{x}], \mathbf{x} \right\rangle.$$

Since $\tilde{\beta} \sim \mathcal{N}(\beta^*, \Psi)$, we have there exists a random vector $\tilde{\theta} \sim \mathcal{N}(\mathbf{0}_d, \mathbf{I}_d)$ such that $\tilde{\beta} = \beta^* + \Psi^{\frac{1}{2}} \tilde{\theta}$ almost surely. Therefore, one has

$$\hat{y}_{\text{Bayes}} = \langle \beta^*, \mathbf{x} \rangle + \left\langle \mathbb{E}[\tilde{\theta} \mid \mathbf{X}, \mathbf{y}, \mathbf{x}], \Psi^{\frac{1}{2}} \mathbf{x} \right\rangle.$$

To compute $\mathbb{E}[\tilde{\theta} \mid \mathbf{X}, \mathbf{y}, \mathbf{x}]$, it suffices to solve the posterior distribution of β given $\mathbf{X}, \mathbf{y}$. From $\tilde{\theta} \sim \mathcal{N}(\mathbf{0}_d, \mathbf{I}_d)$, we have

$$\begin{aligned}\mathbb{P}(\tilde{\theta} \mid \mathbf{X}, \mathbf{y}, \mathbf{x}) &\propto \mathbb{P}(\tilde{\theta}) \mathbb{P}(\mathbf{y} \mid \mathbf{X}, \tilde{\theta}) \propto \exp \left(-\frac{\|\mathbf{y} - \mathbf{X}(\beta^* + \Psi^{\frac{1}{2}} \tilde{\theta})\|_2^2}{2\sigma^2} - \frac{1}{2} \tilde{\theta}^\top \cdot \tilde{\theta} \right) \\ &\propto \exp \left(-\frac{1}{2\sigma^2} \left[\tilde{\theta}^\top \left(\Psi^{\frac{1}{2}} \mathbf{X}^\top \mathbf{X} \Psi^{\frac{1}{2}} + \sigma^2 \mathbf{I}_d \right) \tilde{\theta} - 2(\mathbf{y} - \mathbf{X}\beta^*)^\top \mathbf{X} \Psi^{\frac{1}{2}} \tilde{\theta} \right] \right).\end{aligned}$$

Note that, the function above matches the probability density function of a multivariate Gaussian distribution. Therefore, the posterior mean is given by

$$\mathbb{E}[\tilde{\theta} \mid \mathbf{X}, \mathbf{y}, \mathbf{x}] = \left(\Psi^{\frac{1}{2}} \mathbf{X}^\top \mathbf{X} \Psi^{\frac{1}{2}} + \sigma^2 \mathbf{I}_d \right)^{-1} \Psi^{\frac{1}{2}} \mathbf{X}^\top (\mathbf{y} - \mathbf{X}\beta^*).$$

Therefore, we conclude

$$\hat{y}_{\text{Bayes}} = \mathbf{x}^\top \Psi^{\frac{1}{2}} \left(\Psi^{\frac{1}{2}} \mathbf{X}^\top \mathbf{X} \Psi^{\frac{1}{2}} + \sigma^2 \mathbf{I}_d \right)^{-1} \Psi^{\frac{1}{2}} \mathbf{X}^\top (\mathbf{y} - \mathbf{X}\beta^*) + \mathbf{x}^\top \beta^*. \quad (\text{F.1})$$

□

F.2 Proof of second equation

Proof. Let's first do a variable transformation. We denote

$$\tilde{\mathbf{X}} = \mathbf{X} \Psi^{\frac{1}{2}}, \quad \tilde{\mathbf{x}} = \Psi^{\frac{1}{2}} \mathbf{x}, \quad \tilde{\mathbf{y}} = \mathbf{y} - \mathbf{X}\beta^*, \quad \tilde{\Lambda} = \Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}}. \quad (\text{F.2})$$

Then, we know $\tilde{\mathbf{X}}[i], \tilde{\mathbf{x}} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(\mathbf{0}_d, \tilde{\Lambda})$. We can write the Bayesian optimal estimator in (F.1) as

$$\hat{y}_{\text{Bayes}} = \tilde{\mathbf{x}}^\top \left(\tilde{\mathbf{X}}^\top \tilde{\mathbf{X}} + \sigma^2 \mathbf{I}_d \right)^{-1} \tilde{\mathbf{X}} \tilde{\mathbf{y}} + \mathbf{x}^\top \beta^*.$$

The true label is

$$y = \mathbf{x}^\top \tilde{\beta} + \varepsilon = \mathbf{x}^\top \left(\beta^* + \Psi^{\frac{1}{2}} \tilde{\theta} \right).$$

Therefore, we have

$$\mathcal{L}(\hat{y}_{\text{Bayes}}; \mathbf{X}) - \sigma^2 = \mathbb{E} \left(\tilde{\mathbf{x}}^\top \left(\tilde{\mathbf{X}}^\top \tilde{\mathbf{X}} + \sigma^2 \mathbf{I}_d \right)^{-1} \tilde{\mathbf{X}} \tilde{\mathbf{y}} + \mathbf{x}^\top \beta^* - \mathbf{x}^\top \left(\beta^* + \Psi^{\frac{1}{2}} \tilde{\theta} \right) \right)^2$$

$$= \mathbb{E} \left(\tilde{\mathbf{x}}^\top \left[\left(\tilde{\mathbf{X}}^\top \tilde{\mathbf{X}} + \sigma^2 \mathbf{I}_d \right)^{-1} \tilde{\mathbf{X}} \tilde{\mathbf{y}} - \tilde{\boldsymbol{\theta}} \right] \right)^2 = \mathbb{E} \left\| \hat{\boldsymbol{\theta}} - \tilde{\boldsymbol{\theta}} \right\|_{\Lambda}^2,$$

where

$$\hat{\boldsymbol{\theta}} := \left(\tilde{\mathbf{X}}^\top \tilde{\mathbf{X}} + \sigma^2 \mathbf{I}_d \right)^{-1} \tilde{\mathbf{X}} \tilde{\mathbf{y}}.$$

Note that, this is equivalent to the risk of estimating the ground true linear weight $\tilde{\boldsymbol{\theta}}$ using $\hat{\boldsymbol{\theta}}$ under the Gaussian prior $\tilde{\boldsymbol{\theta}} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$. The estimator $\hat{\boldsymbol{\theta}}$ is a function of transformed input-output pairs in the context $(\tilde{\mathbf{X}}, \tilde{\mathbf{y}})$ and the transformed query input $\mathbf{x}$, and takes the form of standard ridge estimator with regularization coefficient being σ^2 . We then revoke the standard results for the risk of a ridge estimator. Applying the Theorem 1 and 2 in [33], we know for a fixed weight vector $\tilde{\boldsymbol{\theta}}$, with probability at least $1 - \exp(-\Omega(M))$ we have

$$\left\| \hat{\boldsymbol{\theta}} - \tilde{\boldsymbol{\theta}} \right\|_{\mathbf{H}}^2 \simeq \left(\frac{\sigma^2 + \sum_{i>k^*} \phi_i}{M} \right)^2 \left\| \tilde{\boldsymbol{\theta}} \right\|_{\mathbf{H}_{0:k^*}^{-1}}^2 + \left\| \tilde{\boldsymbol{\theta}} \right\|_{\mathbf{H}_{k^*:\infty}}^2 + \sigma^2 \left(\frac{k^*}{M} + \frac{M \sum_{i>k^*} \phi_i^2}{\sigma^2 + \sum_{i>k^*} \phi_i} \right),$$

where $\phi_1 \geq \phi_2 \geq \dots \geq \phi_d \geq 0$ are ordered eigenvalues of Λ .

$$k^* := \min \left\{ k : \phi_k \geq c \cdot \frac{\sigma^2 + \sum_{i>k^*} \phi_i}{M} \right\}$$

and $c > 1$ is an absolute constant. Here, $\mathbf{H}_{0:k^*}$ is SVD approximation with respect to the largest k^* singular values and $\mathbf{H}_{k^*:\infty}$ is the SVD approximation in the remaining singular values. Namely, if we have the eigen-decomposition of $\mathbf{H} = \mathbf{Q} \cdot \text{diag}(\phi_1, \phi_2, \dots, \phi_d) \cdot \mathbf{Q}^\top$, where $\mathbf{Q}$ is an orthogonal matrix, then $\mathbf{H}_{0:k^*}$ and $\mathbf{H}_{k^*:\infty}$ are given by

$$\mathbf{H}_{0:k^*} = \mathbf{Q} \cdot \text{diag}(\phi_1, \phi_2, \dots, \phi_{k^*}, 0, 0, \dots, 0) \cdot \mathbf{Q}^\top, \quad \mathbf{H}_{k^*:\infty} = \mathbf{H} - \mathbf{H}_{0:k^*}.$$

Taking expectation over $\tilde{\boldsymbol{\theta}} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$, we have

$$\begin{aligned} \mathcal{L}(\hat{\mathbf{y}}_{\text{Bayes}}; \mathbf{X}) - \sigma^2 &= \mathbb{E}_{\tilde{\boldsymbol{\theta}} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)} \left\| \hat{\boldsymbol{\theta}}_{\text{Bayes}} - \tilde{\boldsymbol{\theta}} \right\|_{\mathbf{H}}^2 \\ &= \left(\frac{\sigma^2 + \sum_{i>k^*} \phi_i}{M} \right)^2 \cdot \sum_{i \leq k^*} \frac{1}{\phi_i} + \sum_{i > k^*} \phi_i + \sigma^2 \left(\frac{k^*}{M} + \frac{M \sum_{i>k^*} \phi_i^2}{\sigma^2 + \sum_{i>k^*} \phi_i} \right) \end{aligned}$$

Now we simplify this expression. First, we define

$$\bar{\phi} := c \cdot \frac{\sigma^2 + \sum_{i>k^*} \phi_i}{M}.$$

From the definition, we see $\bar{\phi} \geq \frac{c\sigma^2}{M}$. On the other hand, from the assumption that $\text{tr}(\mathbf{H}) = \text{tr}(\Psi^{\frac{1}{2}} \mathbf{H} \Psi^{\frac{1}{2}}) = \sum_{i=1}^d \phi_i \lesssim \sigma^2$, we know that $\bar{\phi} \lesssim \frac{c\sigma^2}{M}$. Combining two parts, we get

$$\bar{\phi} \simeq \frac{\sigma^2}{M}. \quad (\text{F.3})$$

Therefore, we have

$$\begin{aligned} \mathcal{L}(\hat{\mathbf{y}}_{\text{Bayes}}; \mathbf{X}) - \sigma^2 &\simeq \bar{\phi}^2 \cdot \sum_{i \leq k^*} \frac{1}{\phi_i} + \sum_{i > k^*} \phi_i + \frac{\sigma^2}{M} \cdot \left(k^* + \frac{\sum_{i>k^*} \phi_i^2}{\bar{\phi}^2} \right) \\ &\simeq \sum_i \min \left\{ \frac{\bar{\phi}^2}{\phi_i}, \phi_i \right\} + \bar{\phi} \cdot \sum_i \min \left\{ 1, \frac{\phi_i^2}{\bar{\phi}^2} \right\} \quad (\text{from (F.3)}) \\ &\simeq \sum_i \left(\min \left\{ \frac{\bar{\phi}^2}{\phi_i}, \phi_i \right\} + \min \left\{ \bar{\phi}, \frac{\phi_i^2}{\bar{\phi}} \right\} \right) \\ &\simeq \sum_i \min \{ \phi_i, \bar{\phi} \}. \end{aligned}$$

This finishes the proof. $\square$

G Proof of Theorem 6.3

G.1 Training dynamics

Now we consider doing gradient flow on the risk function, or equivalently, on the excess risk function which differs by only a constant:

$$\frac{d\beta}{dt} = -\frac{1}{2} \frac{\partial}{\partial \beta} [\mathcal{R}(\beta, \Gamma) - \min \mathcal{R}(\cdot, \cdot)]; \quad (\text{G.1})$$

$$\frac{d\Gamma}{dt} = -\frac{1}{2} \frac{\partial}{\partial \Gamma} [\mathcal{R}(\beta, \Gamma) - \min \mathcal{R}(\cdot, \cdot)]. \quad (\text{G.2})$$

First, we have the following corollary of Theorem 5.2, which computes the excess risk $\mathcal{R}(\beta, \Gamma) - \min \mathcal{R}(\cdot, \cdot)$.

Corollary G.1 (Excess Risk). *We fix M as the context length. Consider the ICL risk in (3.5), assume the data is generated following Assumption 3.1. Then, we have*

$$\begin{aligned} & \mathcal{R}(\beta, \Gamma) - \min \mathcal{R}(\cdot, \cdot) \\ &= (\beta - \beta^*)^\top \left[(\mathbf{I}_d - \Gamma \mathbf{H})^\top \mathbf{H} (\mathbf{I}_d - \Gamma \mathbf{H}) + \frac{\text{tr}(\mathbf{H} \Gamma^\top \mathbf{H} \Gamma)}{M} \mathbf{H} + \frac{1}{M} \mathbf{H} \Gamma^\top \mathbf{H} \Gamma \mathbf{H} \right] (\beta - \beta^*) \\ &+ \text{tr} \left[\left(\mathbf{H}^{\frac{1}{2}} \Gamma \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \right) \Omega \left(\mathbf{H}^{\frac{1}{2}} \Gamma \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \right)^\top \right]. \end{aligned} \quad (\text{G.3})$$

Proof. This is obtained directly from equation (C.5) and (C.7). $\square$

Now, we write out the differential equations explicitly in the following lemma.

Lemma G.2 (Dynamical system). *The dynamical system of gradient flow described in (G.1) and (G.2) is*

$$\frac{d\beta}{dt} = - \left[(\mathbf{I}_d - \Gamma \mathbf{H})^\top \mathbf{H} (\mathbf{I}_d - \Gamma \mathbf{H}) + \frac{\text{tr}(\mathbf{H} \Gamma^\top \mathbf{H} \Gamma)}{M} \mathbf{H} + \frac{1}{M} \mathbf{H} \Gamma^\top \mathbf{H} \Gamma \mathbf{H} \right] (\beta - \beta^*) \quad (\text{G.4})$$

$$\begin{aligned} \frac{d\Gamma}{dt} &= - \left(\mathbf{H} \Gamma \mathbf{H}^{\frac{1}{2}} \Omega \mathbf{H}^{\frac{1}{2}} - \mathbf{H} \Psi \mathbf{H} \right) - \frac{M+1}{M} \mathbf{H} \Gamma \mathbf{H} (\beta - \beta^*) (\beta - \beta^*)^\top \mathbf{H} \\ &- \frac{1}{M} (\beta - \beta^*)^\top \mathbf{H} (\beta - \beta^*) \cdot \mathbf{H} \Gamma \mathbf{H} + \mathbf{H} (\beta - \beta^*) (\beta - \beta^*)^\top \mathbf{H}. \end{aligned} \quad (\text{G.5})$$

Proof. This can be obtained by directly calculating the derivatives over the excess risk (G.3). To write out the dynamics of β , it suffices to notice that β only attends the first term of the RHS of (G.3), which is a standard quadratic function. For the derivatives of Γ , we first have

$$\begin{aligned} & \frac{1}{2} \frac{\partial}{\partial \Gamma} \text{tr} \left[\left(\mathbf{H}^{\frac{1}{2}} \Gamma \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \right) \Omega \left(\mathbf{H}^{\frac{1}{2}} \Gamma \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \right)^\top \right] \\ &= \mathbf{H}^{\frac{1}{2}} \left(\mathbf{H}^{\frac{1}{2}} \Gamma \mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \right) \Omega^{\frac{1}{2}} \cdot \Omega^{\frac{1}{2}} \mathbf{H}^{\frac{1}{2}} \quad (\text{the sixth equation in Lemma H.3}) \\ &= \mathbf{H} \Gamma \mathbf{H}^{\frac{1}{2}} \Omega \mathbf{H}^{\frac{1}{2}} - \mathbf{H} \Psi \mathbf{H}. \end{aligned}$$

Now it suffices to compute

$$\frac{\partial}{\partial \Gamma} \frac{1}{2} (\beta - \beta^*)^\top \left[(\mathbf{I}_d - \Gamma \mathbf{H})^\top \mathbf{H} (\mathbf{I}_d - \Gamma \mathbf{H}) + \frac{\text{tr}(\mathbf{H} \Gamma^\top \mathbf{H} \Gamma)}{M} \mathbf{H} + \frac{1}{M} \mathbf{H} \Gamma^\top \mathbf{H} \Gamma \mathbf{H} \right] (\beta - \beta^*).$$

Let's compute the partial derivatives separately. From Lemma H.3, we have

$$\begin{aligned}\frac{\partial}{\partial \mathbf{\Gamma}} \frac{1}{2} (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \left(-\mathbf{H} \mathbf{\Gamma}^\top \mathbf{H} \right) (\boldsymbol{\beta} - \boldsymbol{\beta}^*) &= \frac{\partial}{\partial \mathbf{\Gamma}} \frac{1}{2} (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top (-\mathbf{H} \mathbf{\Gamma} \mathbf{H}) (\boldsymbol{\beta} - \boldsymbol{\beta}^*) \\ &= -\frac{1}{2} \mathbf{H} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H}; \\ \frac{\partial}{\partial \mathbf{\Gamma}} \frac{1}{2} \left(\frac{M+1}{M} (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H} \mathbf{\Gamma}^\top \mathbf{H} \mathbf{\Gamma} \mathbf{H} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) \right) &= \frac{M+1}{M} \mathbf{H} \mathbf{\Gamma} \cdot \mathbf{H} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H}; \\ \frac{\partial}{\partial \mathbf{\Gamma}} \frac{1}{2} \left(\frac{\text{tr} \left\{ \mathbf{H} \mathbf{\Gamma}^\top \mathbf{H} \mathbf{\Gamma} \right\}}{M} \cdot (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) \right) &= \frac{1}{M} (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) \cdot \mathbf{H} \mathbf{\Gamma} \mathbf{H}.\end{aligned}$$

Summing over the three equations above and applying the definition of gradient flow, we conclude the dynamics of $\mathbf{\Gamma}$ in (G.5). $\square$

G.2 Proof of the global convergence

Let's now prove Theorem 6.3. Since the first part of Theorem 6.3 is directly implied by the second part, we only deal with the case with general $\mathbf{H}$. In this section, we denote $\lambda_{-1} > 0$ as the minimal non-zero eigenvalue of $\mathbf{H}$. As in the main text, we define

$$\mathcal{H} := \text{Im}(\mathbf{H})$$

and $\mathcal{H}^\perp$ as its orthogonal complement. First, let's prove the convergence of $\boldsymbol{\beta}$.

Lemma G.3 (Convergence of $\boldsymbol{\beta}$). *Under the dynamical system (G.4) and (G.5), one has*

$$\|\mathcal{P}_{\mathcal{H}}(\boldsymbol{\beta}(t)) - \mathcal{P}_{\mathcal{H}}(\boldsymbol{\beta}^*)\|_2^2 \leq \exp\left(\frac{-2\lambda_{-1}t}{M+1}\right) \|\mathcal{P}_{\mathcal{H}}(\boldsymbol{\beta}(0)) - \mathcal{P}_{\mathcal{H}}(\boldsymbol{\beta}^*)\|_2^2, \quad (\text{G.6})$$

which implies

$$\left\| \mathbf{H}^{\frac{1}{2}} (\boldsymbol{\beta}(t) - \boldsymbol{\beta}^*) \right\|_2^2 \leq \lambda_1 \exp\left(\frac{-2\lambda_{-1}t}{M+1}\right) \|\boldsymbol{\beta}(0) - \boldsymbol{\beta}^*\|_2^2, \quad (\text{G.7})$$

and

$$\mathcal{P}_{\mathcal{H}}(\boldsymbol{\beta}(t)) \rightarrow \mathcal{P}_{\mathcal{H}}(\boldsymbol{\beta}^*)$$

when $t \rightarrow \infty$, from arbitrary initialization $\boldsymbol{\beta}(0)$ and $\mathbf{\Gamma}(0)$. Moreover, one has for any $t > 0$, it holds that

$$\mathcal{P}_{\mathcal{H}^\perp}(\boldsymbol{\beta}(t)) = \mathcal{P}_{\mathcal{H}^\perp}(\boldsymbol{\beta}(0)) \quad (\text{G.8})$$

Proof. We first consider the orthogonal projection operator $\mathcal{P}$. From Lemma H.5 and equation (G.4), we know

$$\frac{d\mathcal{P}_{\mathcal{H}}(\boldsymbol{\beta}(t) - \boldsymbol{\beta}^*)}{dt} = -\mathbf{H} \mathbf{H}^+ \mathbf{H}_{\mathbf{\Gamma}} (\boldsymbol{\beta} - \boldsymbol{\beta}^*),$$

where

$$\mathbf{H}_{\mathbf{\Gamma}} := (\mathbf{I}_d - \mathbf{\Gamma} \mathbf{H})^\top \mathbf{H} (\mathbf{I}_d - \mathbf{\Gamma} \mathbf{H}) + \frac{\text{tr}(\mathbf{H} \mathbf{\Gamma}^\top \mathbf{H} \mathbf{\Gamma})}{M} \mathbf{H} + \frac{1}{M} \mathbf{H} \mathbf{\Gamma}^\top \mathbf{H} \mathbf{\Gamma} \mathbf{H}.$$

From the property of pseudo-inverse, we know

$$\frac{d\mathcal{P}_{\mathcal{H}}(\boldsymbol{\beta}(t) - \boldsymbol{\beta}^*)}{dt} = -\mathbf{H}_{\mathbf{\Gamma}} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) = -\mathbf{H}_{\mathbf{\Gamma}} \mathbf{H}^+ \mathbf{H} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) = -\mathbf{H}_{\mathbf{\Gamma}} \mathcal{P}_{\mathcal{H}}(\boldsymbol{\beta} - \boldsymbol{\beta}^*).$$

Therefore, we have

$$\frac{d}{dt} \left[\frac{1}{2} \|\mathcal{P}_{\mathcal{H}}(\boldsymbol{\beta}(t) - \boldsymbol{\beta}^*)\|_2^2 \right] = -\mathcal{P}_{\mathcal{H}}(\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \cdot \mathbf{H}_{\mathbf{\Gamma}} \cdot \mathcal{P}_{\mathcal{H}}(\boldsymbol{\beta} - \boldsymbol{\beta}^*).$$

Notice that

$$\mathbf{H}_{\mathbf{\Gamma}} \succeq \left(\sqrt{\frac{M}{M+1}} \mathbf{I} - \sqrt{\frac{M+1}{M}} \mathbf{\Gamma} \mathbf{H} \right)^\top \mathbf{H} \left(\sqrt{\frac{M}{M+1}} \mathbf{I} - \sqrt{\frac{M+1}{M}} \mathbf{\Gamma} \mathbf{H} \right) + \frac{1}{M+1} \mathbf{H} \succeq \frac{1}{M+1} \mathbf{H}.$$

This implies

$$\begin{aligned} \frac{d}{dt} \left[\frac{1}{2} \|\mathcal{P}_{\mathcal{H}}(\beta(t) - \beta^*)\|_2^2 \right] &\leq -\frac{1}{M+1} \mathcal{P}_{\mathcal{H}}(\beta - \beta^*)^\top \cdot \mathbf{H} \cdot \mathcal{P}_{\mathcal{H}}(\beta - \beta^*) \\ &\leq -\frac{\lambda_{-1}}{M+1} \|\mathcal{P}_{\mathcal{H}}(\beta - \beta^*)\|_2^2, \end{aligned}$$

where the last line comes from Lemma H.5 and λ_{-1} is the minimal non-zero eigenvalue of $\mathbf{H}$. Via standard integration method in ODE, we know this suggests

$$\|\mathcal{P}_{\mathcal{H}}(\beta(t) - \beta^*)\|_2^2 \leq \exp\left(\frac{-2\lambda_{-1}t}{M+1}\right) \|\mathcal{P}_{\mathcal{H}}(\beta(0) - \beta^*)\|_2^2 \rightarrow 0 \quad \text{when } t \rightarrow \infty.$$

This indicates that $\mathcal{P}_{\mathcal{H}}(\beta(t)) \rightarrow \mathcal{P}_{\mathcal{H}}(\beta^*)$ when $t \rightarrow \infty$. Finally, we have

$$\frac{d\mathcal{P}_{\mathcal{H}^\perp}(\beta(t) - \beta^*)}{dt} = -(\mathbf{I}_d - \mathbf{H}\mathbf{H}^\top) \mathbf{H}_\Gamma(\beta - \beta^*) = \mathbf{0}_d,$$

which implies $\mathcal{P}_{\mathcal{H}^\perp}(\beta(t) - \beta^*) = \mathcal{P}_{\mathcal{H}^\perp}(\beta(0) - \beta^*)$ and hence, $\mathcal{P}_{\mathcal{H}^\perp}(\beta(t)) = \mathcal{P}_{\mathcal{H}^\perp}(\beta(0))$ for any $t > 0$. $\square$

Next, let's consider the convergence of Γ . As defined in the main text, we have

$$\mathcal{Z} := \text{Im}(\mathbf{H} \otimes \mathbf{H})$$

and $\mathcal{Z}^\perp$ is its orthogonal complement. We have the following lemma.

Lemma G.4. *Under the dynamical system (G.4) and (G.5), one has*

$$\begin{aligned} \frac{d}{dt} \mathcal{P}_{\mathcal{Z}^\perp}(\Gamma - \Gamma^*) &= \mathbf{0}_{d \times d}, \tag{G.9} \\ \frac{d}{dt} \mathcal{P}_{\mathcal{Z}}(\Gamma - \Gamma^*) &= -\mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}}(\Gamma - \Gamma^*) \cdot \mathbf{H}^{\frac{1}{2}} \Omega \mathbf{H}^{\frac{1}{2}} - \frac{M+1}{M} \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}}(\Gamma - \Gamma^*) \cdot \mathbf{H}(\beta - \beta^*)(\beta - \beta^*)^\top \mathbf{H} \\ &\quad - \frac{1}{M} (\beta - \beta^*)^\top \mathbf{H}(\beta - \beta^*) \cdot \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}}(\Gamma - \Gamma^*) \cdot \mathbf{H} \\ &\quad - \frac{1}{M} (\beta - \beta^*)^\top \mathbf{H}(\beta - \beta^*) \cdot \mathbf{H} \Gamma^* \mathbf{H} - \frac{1}{M} \mathbf{H} \Gamma^* \mathbf{H}(\beta - \beta^*)(\beta - \beta^*)^\top \mathbf{H} \\ &\quad + \frac{1}{M} \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} + \left(\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right) + \sigma^2 \right) \mathbf{I}_d \right) \mathbf{H}^{\frac{1}{2}} (\beta - \beta^*)(\beta - \beta^*)^\top \mathbf{H}. \tag{G.10} \end{aligned}$$

Proof. The first ODE is trivial since the RHS of (G.5) always lies in $\mathcal{Z} := \text{Im}(\mathbf{H} \otimes \mathbf{H})$, so its projection on $\mathcal{Z}^\perp$ always vanishes. To prove the second equation, we can rewrite (G.5) as

$$\begin{aligned} \frac{d\Gamma}{dt} &= - \left(\mathbf{H}(\Gamma - \Gamma^*) \mathbf{H}^{\frac{1}{2}} \Omega \mathbf{H}^{\frac{1}{2}} - \underbrace{\mathbf{H} \Psi \mathbf{H} + \mathbf{H} \Gamma^* \mathbf{H}^{\frac{1}{2}} \Omega \mathbf{H}^{\frac{1}{2}}}_A \right) \\ &\quad - \frac{M+1}{M} \mathbf{H}(\Gamma - \Gamma^*) \mathbf{H}(\beta - \beta^*)(\beta - \beta^*)^\top \mathbf{H} \\ &\quad - \frac{1}{M} (\beta - \beta^*)^\top \mathbf{H}(\beta - \beta^*) \cdot \mathbf{H}(\Gamma - \Gamma^*) \mathbf{H} + \underbrace{\mathbf{H}(\beta - \beta^*)(\beta - \beta^*)^\top \mathbf{H}}_B \\ &\quad - \frac{1}{M} (\beta - \beta^*)^\top \mathbf{H}(\beta - \beta^*) \cdot \mathbf{H} \Gamma^* \mathbf{H} - \underbrace{\frac{M+1}{M} \mathbf{H} \Gamma^* \mathbf{H}(\beta - \beta^*)(\beta - \beta^*)^\top \mathbf{H}}_C. \end{aligned}$$

For term A, recalling $\Gamma^* = \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \mathbf{H}^{-\frac{1}{2}}$, we have

$$\begin{aligned} \mathbf{H} \Gamma^* \mathbf{H}^{\frac{1}{2}} \Omega \mathbf{H}^{\frac{1}{2}} &= \mathbf{H} \Psi \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}} \Omega \mathbf{H}^{\frac{1}{2}} \\ &= \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}} \Omega \mathbf{H}^{\frac{1}{2}} \quad (\Omega \text{ and } \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \text{ commute.}) \end{aligned}$$

$$\begin{aligned}
&= \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega} \mathbf{H}^{\frac{1}{2}} & (\mathbf{H}^{\frac{1}{2}} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}} = \mathbf{H}^{\frac{1}{2}}) \\
&= \mathbf{H} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \boldsymbol{\Omega} \mathbf{H}^{\frac{1}{2}} & (\boldsymbol{\Omega} \text{ and } \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \text{ commute.}) \\
&= \mathbf{H} \boldsymbol{\Psi} \mathbf{H}.
\end{aligned}$$

This suggests $A = 0$. For $B + C$, we have

$$B + C = -\frac{1}{M} \mathbf{H} \boldsymbol{\Gamma}^* \mathbf{H} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H} + (\mathbf{I}_d - \mathbf{H} \boldsymbol{\Gamma}^*) \mathbf{H} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H}.$$

We can compute $(\mathbf{I}_d - \mathbf{H} \boldsymbol{\Gamma}^*) \mathbf{H}$ as

$$\begin{aligned}
(\mathbf{I}_d - \mathbf{H} \boldsymbol{\Gamma}^*) \mathbf{H} &= \mathbf{H}^{\frac{1}{2}} \left(\mathbf{H}^{\frac{1}{2}} - \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \mathbf{H}^{-\frac{1}{2}} \mathbf{H} \right) \\
&= \mathbf{H}^{\frac{1}{2}} \left(\mathbf{H}^{\frac{1}{2}} - \boldsymbol{\Omega}^{-1} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \mathbf{H}^{-\frac{1}{2}} \mathbf{H} \right) & (\boldsymbol{\Omega} \text{ and } \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \text{ commute.}) \\
&= \mathbf{H}^{\frac{1}{2}} \left(\mathbf{H}^{\frac{1}{2}} - \boldsymbol{\Omega}^{-1} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H} \right) & (\mathbf{H}^{\frac{1}{2}} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}} = \mathbf{H}^{\frac{1}{2}}) \\
&= \mathbf{H}^{\frac{1}{2}} \left(\boldsymbol{\Omega}^{-1} \boldsymbol{\Omega} \mathbf{H}^{\frac{1}{2}} - \boldsymbol{\Omega}^{-1} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H} \right) \\
&= \frac{1}{M} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} + \left(\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \right) + \sigma^2 \right) \mathbf{I}_d \right) \mathbf{H}^{\frac{1}{2}}.
\end{aligned}$$

Therefore,

$$\begin{aligned}
B + C &= -\frac{1}{M} \mathbf{H} \boldsymbol{\Gamma}^* \mathbf{H} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H} \\
&\quad + \frac{1}{M} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} + \left(\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \right) + \sigma^2 \right) \mathbf{I}_d \right) \mathbf{H}^{\frac{1}{2}} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H}.
\end{aligned}$$

Bridging $A = 0$ and the result of $B + C$ into the ODE, we have

$$\begin{aligned}
&\frac{d(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)}{dt} \\
&= -\mathbf{H} (\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*) \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega} \mathbf{H}^{\frac{1}{2}} - \frac{M+1}{M} \mathbf{H} (\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*) \mathbf{H} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H} \\
&\quad - \frac{1}{M} (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) \cdot \mathbf{H} (\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*) \mathbf{H} \\
&\quad - \frac{1}{M} (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) \cdot \mathbf{H} \boldsymbol{\Gamma}^* \mathbf{H} - \frac{1}{M} \mathbf{H} \boldsymbol{\Gamma}^* \mathbf{H} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H} \\
&\quad + \frac{1}{M} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} + \left(\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \right) + \sigma^2 \right) \mathbf{I}_d \right) \mathbf{H}^{\frac{1}{2}} (\boldsymbol{\beta} - \boldsymbol{\beta}^*) (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H}. \quad (\text{G.11})
\end{aligned}$$

Note that, all terms in the RHS above lie in $\mathcal{Z} = \text{Im}(\mathbf{H}^{\otimes 2})$, so this summation also equals $\frac{d}{dt} \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)$. Moreover, since

$$\begin{aligned}
\mathbf{H}^{\frac{1}{2}} (\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*) \mathbf{H}^{\frac{1}{2}} &= \mathbf{H}^{\frac{1}{2}} \left[\mathcal{P}_{\text{Im}(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}})} (\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*) + \mathcal{P}_{\text{null}(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}})} (\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*) \right] \mathbf{H}^{\frac{1}{2}} \\
&= \mathbf{H}^{\frac{1}{2}} \cdot \mathcal{P}_{\text{Im}(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}})} (\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*) \mathbf{H}^{\frac{1}{2}}. \quad (\text{G.12})
\end{aligned}$$

Bridging (G.12) into (G.11), we conclude. $\square$

Then, let's upper bound the dynamics of $\|\mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)\|_F^2$ in the following lemma.

Lemma G.5 (Dynamics of Frobenius norm). *Under the dynamical system (G.4) and (G.5), one has*

$$\frac{d}{dt} \left[\frac{1}{2} \|\mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)\|_F^2 \right] \leq -A_1 \|\mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)\|_F^2 + A_2 \exp \left(\frac{-2\lambda_{-1}t}{M+1} \right) \|\mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)\|_F, \quad (\text{G.13})$$

where

$$\begin{aligned} A_1 &= \lambda_{-1}^2 \lambda_{\min}(\boldsymbol{\Omega}), \\ A_2 &= \left(2 + \frac{2}{M}\right) \lambda_1^2 \|\boldsymbol{\beta}(0) - \boldsymbol{\beta}^*\|_2^2. \end{aligned} \quad (\text{G.14})$$

are two positive constant. Here, $\lambda_{-1} > 0$ is the minimal positive eigenvalue of $\mathbf{H}$, λ_1 is the maximal eigenvalue of $\mathbf{H}$, $\lambda_{\min}(\boldsymbol{\Omega})$ is the minimal eigenvalue of $\boldsymbol{\Omega}$ (defined in (A.3)) and is strictly positive.

Proof. From the dynamics of $\mathcal{P}_{\mathcal{Z}}\boldsymbol{\Gamma}$ in (G.10), we know

$$\begin{aligned} \frac{d}{dt} \left[\frac{1}{2} \|\mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)\|_F^2 \right] &= \text{tr} \left(\frac{d}{dt} \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*) \cdot \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)^\top \right) \\ &= G_1 + G_2 + G_3, \end{aligned} \quad (\text{G.15})$$

where

$$\begin{aligned} G_1 &= -\text{tr} \left[\mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*) \cdot \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega} \mathbf{H}^{\frac{1}{2}} \cdot \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)^\top \right], \\ G_2 &= -\text{tr} \left[\frac{M+1}{M} \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*) \cdot \mathbf{H}(\boldsymbol{\beta} - \boldsymbol{\beta}^*)(\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)^\top \right] \\ &\quad - \text{tr} \left[\frac{1}{M} (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H}(\boldsymbol{\beta} - \boldsymbol{\beta}^*) \cdot \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*) \cdot \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)^\top \right], \\ G_3 &= -\text{tr} \left[\frac{1}{M} (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H}(\boldsymbol{\beta} - \boldsymbol{\beta}^*) \cdot \mathbf{H} \boldsymbol{\Gamma}^* \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)^\top \right] \\ &\quad + \frac{1}{M} \text{tr} \left[\mathbf{H} \boldsymbol{\Gamma}^* \mathbf{H}(\boldsymbol{\beta} - \boldsymbol{\beta}^*)(\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)^\top \right] \\ &\quad + \text{tr} \left[\frac{1}{M} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} + \left(\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \right) + \sigma^2 \right) \mathbf{I}_d \right) \mathbf{H}^{\frac{1}{2}} (\boldsymbol{\beta} - \boldsymbol{\beta}^*)(\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)^\top \right]. \end{aligned} \quad (\text{G.16})$$

From Lemma H.2, we know that $G_2 \leq 0$. Then, we consider G_1 and G_3 . For G_1 , we have

$$\begin{aligned} G_1 &= -\text{tr} \left[\left(\mathbf{H}^{\frac{1}{2}} \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)^\top \mathbf{H}^{\frac{1}{2}} \right) \cdot \left(\mathbf{H}^{\frac{1}{2}} \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*) \mathbf{H}^{\frac{1}{2}} \right) \cdot \boldsymbol{\Omega} \right] \\ &\leq -\lambda_{-1}^2 \text{tr} \left[\mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)^\top \cdot \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*) \cdot \boldsymbol{\Omega} \right] \quad ((3) \text{ in Lemma H.6 and } (2) \text{ in Lemma H.2}) \\ &\leq -\lambda_{-1}^2 \lambda_{\min}(\boldsymbol{\Omega}) \|\mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)\|_F^2, \end{aligned}$$

where $\lambda_{\min}(\cdot)$ denote the minimal eigenvalue. Since $\boldsymbol{\Omega}$ is positive definite, we know $\lambda_{\min}(\boldsymbol{\Omega}) > 0$.

Finally, let's upper bound G_3 . First, we have

$$\begin{aligned} & -\text{tr} \left[\frac{1}{M} (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H}(\boldsymbol{\beta} - \boldsymbol{\beta}^*) \cdot \mathbf{H} \boldsymbol{\Gamma}^* \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)^\top \right] \\ & \leq \frac{\sqrt{d}}{M} (\boldsymbol{\beta} - \boldsymbol{\beta}^*)^\top \mathbf{H}(\boldsymbol{\beta} - \boldsymbol{\beta}^*) \cdot \left\| \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Gamma}^* \mathbf{H}^{\frac{1}{2}} \right\|_2 \left\| \mathbf{H}^{\frac{1}{2}} \cdot \mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)^\top \mathbf{H}^{\frac{1}{2}} \right\|_F \quad (\text{Lemma H.1}) \\ & \leq \frac{\sqrt{d} \lambda_1}{M} \exp \left(\frac{-2\lambda_{-1}t}{M+1} \right) \|\boldsymbol{\beta}(0) - \boldsymbol{\beta}^*\|_2^2 \lambda_{\max} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Gamma}^* \mathbf{H}^{\frac{1}{2}} \right) \cdot \|\mathbf{H}\|_F \|\mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)\|_F \\ & \quad (\text{from equation (G.7)}) \\ & \leq \frac{\sqrt{d} \lambda_1^2}{M} \exp \left(\frac{-2\lambda_{-1}t}{M+1} \right) \|\boldsymbol{\beta}(0) - \boldsymbol{\beta}^*\|_2^2 \lambda_{\max} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Gamma}^* \mathbf{H}^{\frac{1}{2}} \right) \cdot \|\mathcal{P}_{\mathcal{Z}}(\boldsymbol{\Gamma} - \boldsymbol{\Gamma}^*)\|_F. \end{aligned}$$

For $\lambda_{\max} \left(\mathbf{H}^{\frac{1}{2}} \boldsymbol{\Gamma}^* \mathbf{H}^{\frac{1}{2}} \right)$, we recall $\boldsymbol{\Gamma}^* = \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \mathbf{H}^{-\frac{1}{2}}$ and obtain

$$\begin{aligned} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Gamma}^* \mathbf{H}^{\frac{1}{2}} &= \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}} = \boldsymbol{\Omega}^{-1} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}} \quad (\boldsymbol{\Omega} \text{ and } \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \text{ commute}) \\ &= \boldsymbol{\Omega}^{-1} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}}. \quad (\mathbf{H}^{\frac{1}{2}} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}} = \mathbf{H}^{\frac{1}{2}}) \end{aligned}$$

Since Ω^{-1} and $\mathbf{H}^{\frac{1}{2}} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}}$ commute, they are simultaneously diagonalizable. The eigenvalues of $\Omega^{-1} \mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}}$ are

$$\frac{\phi_i}{\frac{M+1}{M} \phi_i + \frac{\sum_i \phi_i + \sigma^2}{M}}, \quad i = 1, 2, \dots, d.$$

Note that, every eigenvalue is upper bounded by 1, so we simply get

$$\lambda_{\max} \left(\mathbf{H}^{\frac{1}{2}} \Gamma^* \mathbf{H}^{\frac{1}{2}} \right) \leq 1. \quad (\text{G.17})$$

Therefore, we have

$$\begin{aligned} & -\text{tr} \left[\frac{1}{M} (\beta - \beta^*)^\top \mathbf{H} (\beta - \beta^*) \cdot \mathbf{H} \Gamma^* \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)^\top \right] \\ & \leq \frac{\sqrt{d} \lambda_1^2}{M} \exp \left(\frac{-2\lambda_{-1}t}{M+1} \right) \|\beta(0) - \beta^*\|_2^2 \cdot \|\mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)\|_F. \end{aligned} \quad (\text{G.18})$$

Similarly, we have

$$\begin{aligned} & -\text{tr} \left[\frac{1}{M} \mathbf{H} \Gamma^* \mathbf{H} (\beta - \beta^*) (\beta - \beta^*)^\top \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)^\top \right] \\ & \leq \frac{\sqrt{d}}{M} \left\| \mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)^\top \right\|_F \left\| \mathbf{H}^{\frac{1}{2}} (\beta - \beta^*) (\beta - \beta^*)^\top \mathbf{H}^{\frac{1}{2}} \right\|_F \|\mathbf{H}\|_F \left\| \mathbf{H}^{\frac{1}{2}} \Gamma^* \mathbf{H}^{\frac{1}{2}} \right\|_2 \quad (\text{Lemma H.1}) \\ & \leq \frac{\sqrt{d} \lambda_1^2}{M} \exp \left(\frac{-2\lambda_{-1}t}{M+1} \right) \|\beta(0) - \beta^*\|_2^2 \cdot \|\mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)\|_F, \end{aligned} \quad (\text{G.19})$$

and

$$\begin{aligned} & \text{tr} \left[\frac{1}{M} \mathbf{H}^{\frac{1}{2}} \Omega^{-1} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} + \left(\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right) + \sigma^2 \right) \mathbf{I}_d \right) \mathbf{H}^{\frac{1}{2}} (\beta - \beta^*) (\beta - \beta^*)^\top \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)^\top \right] \\ & \leq \frac{\sqrt{d}}{M} \left\| \Omega^{-1} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} + \left(\text{tr} \left(\mathbf{H}^{\frac{1}{2}} \Psi \mathbf{H}^{\frac{1}{2}} \right) + \sigma^2 \right) \mathbf{I}_d \right) \right\|_2 \left\| \mathbf{H}^{\frac{1}{2}} (\beta - \beta^*) (\beta - \beta^*)^\top \mathbf{H} \cdot \mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)^\top \mathbf{H}^{\frac{1}{2}} \right\|_F \\ & \quad (\text{Lemma H.1}) \\ & \leq \frac{\sqrt{d}}{M} \cdot M \cdot \left\| \mathbf{H}^{\frac{1}{2}} (\beta - \beta^*) (\beta - \beta^*)^\top \mathbf{H}^{\frac{1}{2}} \right\|_F \cdot \|\mathbf{H}\|_F \cdot \|\mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)\|_F \\ & \leq \sqrt{d} \lambda_1^2 \exp \left(\frac{-2\lambda_{-1}t}{M+1} \right) \|\beta(0) - \beta^*\|_2^2 \cdot \|\mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)\|_F \end{aligned} \quad (\text{G.20})$$

Bridging (G.18), (G.19) and (G.20) into (G.16), we get

$$G_3 \leq \left(2 + \frac{2}{M} \right) \lambda_1^2 \exp \left(\frac{-2\lambda_{-1}t}{M+1} \right) \|\beta(0) - \beta^*\|_2^2 \cdot \|\mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)\|_F \quad (\text{G.21})$$

□

Finally, we finish this section by proving the convergence of Γ .

Lemma G.6 (Convergence of Γ). *Under the dynamical system (G.4) and (G.5), one has*

$$\mathcal{P}_{\mathcal{Z}} (\Gamma(t)) \rightarrow \mathcal{P}_{\mathcal{Z}} (\Gamma^*), \quad \mathcal{P}_{\mathcal{Z}^\perp} (\Gamma(t)) = \mathcal{P}_{\mathcal{Z}^\perp} (\Gamma(0))$$

when $t \rightarrow \infty$, from arbitrary initialization $\beta(0)$ and $\Gamma(0)$.

Proof. We observe that

$$\frac{d}{dt} \left[\frac{1}{2} \|\mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)\|_F^2 \right] = \|\mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)\|_F \cdot \frac{d}{dt} \|\mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)\|_F.$$

Combining it with Lemma G.5, we have

$$\frac{d}{dt} \|\mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)\|_F \leq -A_1 \|\mathcal{P}_{\mathcal{Z}} (\Gamma - \Gamma^*)\|_F + A_2 \exp \left(\frac{-2\lambda_{-1}t}{M+1} \right),$$

where $A_1 > 0, A_2 > 0$ are defined in (G.14). Simple calculation shows that

$$\frac{d}{dt} [\exp(A_1 t) \|\mathcal{P}_{\mathcal{Z}}(\Gamma - \Gamma^*)\|_F] \leq A_2 \exp \left[\left(A_1 - \frac{2\lambda_{-1}}{M+1} \right) t \right]. \quad (\text{G.22})$$

When $A_1 \neq \frac{2\lambda_{-1}}{M+1}$, we integrate both sides from $t = 0$ to $t = T$, then divide them by $\exp(A_1 T)$. This gives

$$\begin{aligned} \|\mathcal{P}_{\mathcal{Z}}(\Gamma(T) - \Gamma^*)\|_F &\leq \frac{\|\mathcal{P}_{\mathcal{Z}}(\Gamma(0) - \Gamma^*)\|_F}{\exp(A_1 T)} + \frac{A_2}{A_1 - \frac{2\lambda_{-1}}{M+1}} \cdot \frac{\exp \left[\left(A_1 - \frac{2\lambda_{-1}}{M+1} \right) T \right] - 1}{\exp(A_1 T)} \\ &= \frac{\|\mathcal{P}_{\mathcal{Z}}(\Gamma(0) - \Gamma^*)\|_F}{\exp(A_1 T)} + \frac{A_2}{A_1 - \frac{2\lambda_{-1}}{M+1}} \cdot \left[\exp \left(-\frac{2\lambda_{-1}T}{M+1} \right) - \exp(-A_1 T) \right] \rightarrow 0 \end{aligned}$$

when $T \rightarrow \infty$. Otherwise, if $A_1 = \frac{2\lambda_{-1}}{M+1}$, the right hand side of (G.22) reduces to a constant, which implies $\exp(A_1 t) \|\mathcal{P}_{\mathcal{Z}}(\Gamma(t) - \Gamma^*)\|_F$ grows at most at a linear rate, which implies $\|\mathcal{P}_{\mathcal{Z}}(\Gamma(t) - \Gamma^*)\|_F \rightarrow 0$ when $t \rightarrow \infty$. Together, in all cases, we have $\|\mathcal{P}_{\mathcal{Z}}(\Gamma(t) - \Gamma^*)\|_F \rightarrow 0$. From (G.9), we soon get $\mathcal{P}_{\mathcal{Z}}(\Gamma(t)) = \mathcal{P}_{\mathcal{Z}}(\Gamma(0))$ for all $t > 0$. Therefore, we conclude. $\square$

The Theorem 6.3 is proved by simply combining Lemma G.3 and Lemma G.6.

H Technical lemmas

Lemma H.1 (Von-Neumann's Trace Inequality). *Let $U, V \in \mathbb{R}^{d \times n}$ with $d \leq n$. We have*

$$\operatorname{tr}(U^\top V) \leq \sum_{i=1}^d \sigma_i(U) \sigma_i(V) \leq \|U\|_{\text{op}} \times \sum_{i=1}^d \sigma_i(V) \leq \sqrt{d} \cdot \|U\|_{\text{op}} \|V\|_F$$

where $\sigma_1(X) \geq \sigma_2(X) \geq \dots \geq \sigma_d(X)$ are the ordered singular values of $X \in \mathbb{R}^{d \times n}$.

Lemma H.2 ([22]). *For any two positive semi-definite matrices $A, B \in \mathbb{R}^{d \times d}$, we have*

- $\operatorname{tr}[AB] \geq 0$.
- $AB \succeq 0$ if and only if A and B commute.

Lemma H.3 (Derivatives, [27]). *We denote $\mathbf{A}, \mathbf{B}, \mathbf{C}, \mathbf{D}, \mathbf{X}$ as matrices and $\mathbf{a}, \mathbf{b}, \mathbf{x}$ as vectors. Then, we have*

- $\frac{\partial}{\partial \mathbf{X}} \operatorname{tr}[\mathbf{X}^\top \mathbf{B} \mathbf{X} \mathbf{C}] = \mathbf{B} \mathbf{X} \mathbf{C} + \mathbf{B}^\top \mathbf{X} \mathbf{C}^\top$.
- $\frac{\partial \mathbf{x}^\top \mathbf{B} \mathbf{x}}{\partial \mathbf{x}} = (\mathbf{B} + \mathbf{B}^\top) \mathbf{x}$.
- $\frac{\partial \mathbf{a}^\top \mathbf{X} \mathbf{b}}{\partial \mathbf{X}} = \mathbf{a} \mathbf{b}^\top$.
- $\frac{\partial \mathbf{a}^\top \mathbf{X}^\top \mathbf{b}}{\partial \mathbf{X}} = \mathbf{b} \mathbf{a}^\top$.
- $\frac{\partial \mathbf{b}^\top \mathbf{X}^\top \mathbf{D} \mathbf{X} \mathbf{c}}{\partial \mathbf{X}} = \mathbf{D}^\top \mathbf{X} \mathbf{b} \mathbf{c}^\top + \mathbf{D} \mathbf{X} \mathbf{c} \mathbf{b}^\top$.
- $\frac{\partial}{\partial \mathbf{X}} \operatorname{tr}[(\mathbf{A} \mathbf{X} \mathbf{B} + \mathbf{C})(\mathbf{A} \mathbf{X} \mathbf{B} + \mathbf{C})^\top] = 2\mathbf{A}^\top (\mathbf{A} \mathbf{X} \mathbf{B} + \mathbf{C}) \mathbf{B}^\top$

Lemma H.4 (Lemma D.2 in [38], Lemma 4.2 in [37]). *If $\mathbf{x}$ is Gaussian random vector of d dimension, mean zero and covariance matrix $\mathbf{H}$, and $\mathbf{A} \in \mathbb{R}^{d \times d}$ is a fixed matrix. Then*

$$\mathbb{E}[\mathbf{x} \mathbf{x}^\top \mathbf{A} \mathbf{x} \mathbf{x}^\top] = \mathbf{H} (\mathbf{A} + \mathbf{A}^\top) \mathbf{H} + \operatorname{tr}(\mathbf{A} \mathbf{H}) \mathbf{H}.$$

If $\mathbf{A}$ is symmetric and the rows in $\mathbf{X} \in \mathbb{R}^{M \times d}$ are generated independently from

$$\mathbf{X}[i] \sim \mathcal{N}(0, \mathbf{H}), \quad i = 1, \dots, M.$$

Then, it holds that

$$\mathbb{E}[\mathbf{X}^\top \mathbf{X} \mathbf{A} \mathbf{X}^\top \mathbf{X}] = M \cdot \operatorname{tr}(\mathbf{H} \mathbf{A}) \cdot \mathbf{H} + M(M+1) \cdot \mathbf{H} \mathbf{A} \mathbf{H}.$$

Lemma H.5 (Linear Algebra). *Suppose $\mathbf{H}$ is a (non-zero) positive semi-definite matrix in $\mathbb{R}^{d \times d}$ and $\mathbf{H}^{\frac{1}{2}}$ is its principle square root. We denote λ_{-1} as the minimal non-zero eigenvector of $\mathbf{H}$. Then, we have*

- 1. $\operatorname{null}(\mathbf{H}) = \operatorname{null}(\mathbf{H}^{\frac{1}{2}}), \operatorname{Im}(\mathbf{H}) = \operatorname{Im}(\mathbf{H}^{\frac{1}{2}})$.
- 2. $\mathbf{H} \mathbf{H}^+ = \mathbf{H}^+ \mathbf{H}$, where $(\cdot)^+$ denotes the Moore-Penrose pseudo-inverse.
- 3. For any vector $\boldsymbol{\alpha} \in \mathbb{R}^d$, the orthogonal projection operator on $\operatorname{Im}(\mathbf{H})$ and $\operatorname{null}(\mathbf{H})$ are respectively

$$\mathcal{P}_{\operatorname{Im}(\mathbf{H})}(\boldsymbol{\alpha}) = \mathbf{H} \mathbf{H}^+ \boldsymbol{\alpha} = \mathbf{H}^+ \mathbf{H} \boldsymbol{\alpha}, \quad \mathcal{P}_{\operatorname{null}(\mathbf{H})}(\boldsymbol{\alpha}) = (\mathbf{I}_d - \mathbf{H} \mathbf{H}^+) \boldsymbol{\alpha} = (\mathbf{I}_d - \mathbf{H}^+ \mathbf{H}) \boldsymbol{\alpha}.$$
- 4. For any vector $\boldsymbol{\alpha} \in \mathbb{R}^d$, we have

$$\mathcal{P}_{\operatorname{Im}(\mathbf{H})}(\boldsymbol{\alpha})^\top \mathbf{H} \mathcal{P}_{\operatorname{Im}(\mathbf{H})}(\boldsymbol{\alpha}) \geq \lambda_{-1} \|\mathcal{P}_{\operatorname{Im}(\mathbf{H})}(\boldsymbol{\alpha})\|_2^2$$

Proof. We consider the eigen-decomposition of matrix $\mathbf{H}$, which is

$$\mathbf{H} = \mathbf{Q} \mathbf{D} \mathbf{Q}^\top, \quad (\text{H.1})$$

where $\mathbf{D} = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_d)$ is a diagonal matrix with diagonal entries being the eigenvalues of $\mathbf{H}$, and $\mathbf{Q}$ is an orthogonal matrix. Then, from the definition of principle square root, we know

$$\mathbf{H}^{\frac{1}{2}} = \mathbf{Q} \mathbf{D}^{\frac{1}{2}} \mathbf{Q}^\top, \quad (\text{H.2})$$

where $\mathbf{D}^{\frac{1}{2}} = \text{diag}(\sqrt{\lambda_1}, \sqrt{\lambda_2}, \dots, \sqrt{\lambda_d})$. We denote columns of $\mathbf{Q}$ as $\mathbf{q}_1, \dots, \mathbf{q}_d$. We know they are the eigenvectors of $\mathbf{H}$. From the eigen-decomposition of $\mathbf{H}$ and $\mathbf{H}^{\frac{1}{2}}$, we know they share the same set of eigenvectors. Without loss of generality, we assume $\text{rank}(\mathbf{H}) = r$ and $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_r > 0, \lambda_i = 0, i = r+1, \dots, d$. Then, we denote $\mathcal{H}$ as the linear vector space spanned by $\mathbf{q}_1, \dots, \mathbf{q}_r$ and $\mathcal{H}^\perp$ as its orthogonal complement (which is the subspace spanned by $\mathbf{q}_{r+1}, \dots, \mathbf{q}_d$). For any vector $\boldsymbol{\alpha} \in \mathbb{R}^d$, we can write its orthogonal decomposition as $\boldsymbol{\alpha} = \sum_{i=1}^d a_i \mathbf{q}_i$, so we have

$$\mathbf{H} \boldsymbol{\alpha} = \sum_{i=1}^d a_i \mathbf{H} \mathbf{q}_i = \sum_{i=1}^r a_i \lambda_i \mathbf{q}_i \in \mathcal{H},$$

which implies $\text{Im}(\mathbf{H}) \subset \mathcal{H}$. On the other hand, we know for any $\boldsymbol{\alpha} \in \mathcal{H}$, we can write it as $\boldsymbol{\alpha} = \sum_{i=1}^r a_i \mathbf{q}_i$ for some $a_1, \dots, a_r$, then we have

$$\boldsymbol{\alpha} = \sum_{i=1}^r \frac{a_i}{\lambda_i} \cdot \lambda_i \mathbf{q}_i = \mathbf{H} \left(\sum_{i=1}^r \frac{a_i}{\lambda_i} \mathbf{q}_i \right) \in \text{Im}(\mathbf{H}),$$

which implies $\mathcal{H} \subset \text{Im}(\mathbf{H})$. This shows

$$\text{Im}(\mathbf{H}) = \mathcal{H} = \text{span}\{\mathbf{q}_1, \dots, \mathbf{q}_r\},$$

and

$$\text{null}(\mathbf{H}) = \mathcal{H}^\perp = \text{span}\{\mathbf{q}_{r+1}, \dots, \mathbf{q}_d\},$$

since $\text{null}(\mathbf{H})$ is the orthogonal complement of $\text{Im}(\mathbf{H})$. Then, (1) is proved by noticing that $\mathbf{H}$ and $\mathbf{H}^{\frac{1}{2}}$ share the same set of eigenvectors. To prove (2), it suffices to notice that

$$\mathbf{H} \mathbf{H}^+ = \mathbf{Q} \cdot \underbrace{\text{diag}(1, 1, \dots, 1, 0, 0, \dots, 0)}_r \cdot \mathbf{Q}^\top = \mathbf{H}^+ \mathbf{H}.$$

To prove (3), it suffices to notice that actually $\mathbf{H} \mathbf{H}^+ \boldsymbol{\alpha} \in \text{Im}(\mathbf{H})$ and

$$\langle \mathbf{H} \mathbf{H}^+ \boldsymbol{\alpha}, (\mathbf{I}_d - \mathbf{H} \mathbf{H}^+) \boldsymbol{\alpha} \rangle = \langle \mathbf{H}^+ \mathbf{H} \boldsymbol{\alpha}, (\mathbf{I}_d - \mathbf{H} \mathbf{H}^+) \boldsymbol{\alpha} \rangle = \boldsymbol{\alpha}^\top (\mathbf{H} \mathbf{H}^+ - \mathbf{H} \mathbf{H}^+ \mathbf{H} \mathbf{H}^+) \boldsymbol{\alpha} = 0.$$

Finally, to prove (4), we can write $\boldsymbol{\alpha} = \sum_{i=1}^d a_i \mathbf{q}_i$ for some $a_1, \dots, a_d$. Then, by orthogonality we have

$$\begin{aligned} \mathcal{P}_{\text{Im}(\mathbf{H})}(\boldsymbol{\alpha})^\top \mathbf{H} \mathcal{P}_{\text{Im}(\mathbf{H})}(\boldsymbol{\alpha}) &= \left(\sum_{i=1}^d a_i \mathbf{q}_i \right)^\top \mathbf{H} \left(\sum_{i=1}^d a_i \mathbf{q}_i \right) = \sum_{i=1}^r a_i^2 \lambda_i \mathbf{q}_i^\top \mathbf{q}_i \geq \lambda_{-1} \cdot \sum_{i=1}^r a_i^2 \mathbf{q}_i^\top \mathbf{q}_i \\ &= \lambda_{-1} \|\mathcal{P}_{\text{Im}(\mathbf{H})}(\boldsymbol{\alpha})\|_2^2. \end{aligned}$$

Therefore, we conclude. $\square$

Lemma H.6 (Tensor product). Suppose $\mathbf{H}$ is a (non-zero) positive semi-definite matrix in $\mathbb{R}^{d \times d}$ and $\mathbf{H}^{\frac{1}{2}}$ is its principle square root. We denote $\otimes$ as Kronecker product, which is defined as

$$(\mathbf{A} \otimes \mathbf{B}) \circ \mathbf{C} = \mathbf{B} \mathbf{C} \mathbf{A}^\top.$$

We define

$$\text{Im}(\mathbf{H} \otimes \mathbf{H}) := \{(\mathbf{H} \otimes \mathbf{H}) \circ \mathbf{Z} : \mathbf{Z} \in \mathbb{R}^{d \times d}\}, \quad \text{null}(\mathbf{H} \otimes \mathbf{H}) := \{\mathbf{Z} \in \mathbb{R}^{d \times d} : (\mathbf{H} \otimes \mathbf{H}) \circ \mathbf{Z} = \mathbf{0}\},$$

and define $\text{Im}(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}})$ and $\text{null}(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}})$ similarly. Then, we have

- 1. $\text{Im}(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}}) = \text{Im}(\mathbf{H} \otimes \mathbf{H})$, $\text{null}(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}}) = \text{null}(\mathbf{H} \otimes \mathbf{H})$.
- 2. We denote $\mathcal{Z} = \text{Im}(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}}) = \text{Im}(\mathbf{H} \otimes \mathbf{H})$ and $\mathcal{Z}^\perp = \text{null}(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}}) = \text{null}(\mathbf{H} \otimes \mathbf{H})$. For any matrix $\mathbf{Z} \in \mathbb{R}^{d \times d}$, we have

$$\mathcal{P}_{\mathcal{Z}}(\mathbf{Z}) = (\mathbf{H}\mathbf{H}^+)^{\otimes 2} \circ \mathbf{Z} = \mathbf{H}\mathbf{H}^+ \mathbf{Z} \mathbf{H}^+ \mathbf{H} = \mathbf{H}^{\frac{1}{2}} \mathbf{H}^{-\frac{1}{2}} \mathbf{Z} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}} \quad (\text{H.3})$$

$$\mathcal{P}_{\mathcal{Z}^\perp}(\mathbf{Z}) = [\mathbf{I}_d^{\otimes 2} - (\mathbf{H}\mathbf{H}^+)^{\otimes 2}] \circ \mathbf{Z} = \mathbf{Z} - \mathbf{H}\mathbf{H}^+ \mathbf{Z} \mathbf{H}^+ \mathbf{H} = \mathbf{Z} - \mathbf{H}^{\frac{1}{2}} \mathbf{H}^{-\frac{1}{2}} \mathbf{Z} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}}. \quad (\text{H.4})$$

- 3. For any matrix $\mathbf{Z} \in \mathbb{R}^{d \times d}$, it holds that

$$\left((\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}}) \circ \mathcal{P}_{\mathcal{Z}}(\mathbf{Z}) \right) \cdot \left((\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}}) \circ \mathcal{P}_{\mathcal{Z}}(\mathbf{Z}) \right)^\top \succeq \lambda_{-1}^2 \mathcal{P}_{\mathcal{Z}}(\mathbf{Z}) \cdot \mathcal{P}_{\mathcal{Z}}(\mathbf{Z})^\top, \quad (\text{H.5})$$

where λ_{-1} is the minimal positive eigenvector of $\mathbf{H}$.

Proof. Let's first prove the second part. From the definition of tensor product and the fact that $\mathbf{H}\mathbf{H}^+ = \mathbf{H}^+ \mathbf{H}$ (see Lemma H.5), we know the second equation in H.3 holds. To prove the first equation, it suffices to notice that for any matrix $\mathbf{Z} \in \mathbb{R}^{d \times d}$, the matrix $\mathbf{H}\mathbf{H}^+ \mathbf{Z} \mathbf{H}^+ \mathbf{H}$ is in $\text{Im}(\mathbf{H}^{\otimes 2})$, together with the fact that

$$\begin{aligned} \langle \mathbf{H}\mathbf{H}^+ \mathbf{Z} \mathbf{H}^+ \mathbf{H}, \mathbf{Z} - \mathbf{H}\mathbf{H}^+ \mathbf{Z} \mathbf{H}^+ \mathbf{H} \rangle &= \text{tr}(\mathbf{H}\mathbf{H}^+ \mathbf{Z} \mathbf{H}^+ \mathbf{H} \mathbf{Z}^\top) - \text{tr}(\mathbf{H}\mathbf{H}^+ \mathbf{Z} \mathbf{H}^+ \mathbf{H} \mathbf{H} \mathbf{H}^+ \mathbf{Z}^\top \mathbf{H}^+ \mathbf{H}) \\ &= \text{tr}(\mathbf{H}\mathbf{H}^+ \mathbf{Z} \mathbf{H}^+ \mathbf{H} \mathbf{Z}^\top) - \text{tr}(\mathbf{H}\mathbf{H}^+ \mathbf{H} \mathbf{H}^+ \mathbf{Z} \mathbf{H}^+ \mathbf{H} \mathbf{H}^+ \mathbf{H} \mathbf{Z}^\top) \\ &\quad (\mathbf{H}\mathbf{H}^+ = \mathbf{H}^+ \mathbf{H}) \\ &= 0. \end{aligned}$$

The proof of (H.4) is similar. Then, from Lemma H.5, we know $\text{Im}(\mathbf{H}) = \text{Im}(\mathbf{H}^{\frac{1}{2}})$, which implies the projection operator onto those two subspace are identical, indicating $\mathbf{H}\mathbf{H}^+ = \mathbf{H}^{\frac{1}{2}} \mathbf{H}^{-\frac{1}{2}}$. Therefore, we have

$$\mathcal{P}_{\text{Im}(\mathbf{H} \otimes \mathbf{H})}(\mathbf{Z}) = (\mathbf{H}\mathbf{H}^+)^{\otimes 2} \circ \mathbf{Z} = \mathbf{H}\mathbf{H}^+ \mathbf{Z} \mathbf{H}^+ \mathbf{H} = \mathbf{H}^{\frac{1}{2}} \mathbf{H}^{-\frac{1}{2}} \mathbf{Z} \mathbf{H}^{-\frac{1}{2}} \mathbf{H}^{\frac{1}{2}} = \mathcal{P}_{\text{Im}(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}})}(\mathbf{Z}).$$

Since this holds for any matrix $\mathbf{Z} \in \mathbb{R}^{d \times d}$, this suggests $\text{Im}(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}}) = \text{Im}(\mathbf{H} \otimes \mathbf{H})$. Similarly, we also have $\text{null}(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}}) = \text{null}(\mathbf{H} \otimes \mathbf{H})$. Finally, to prove (3), we use the eigen-decomposition of $\mathbf{H}$ and $\mathbf{H}^{\frac{1}{2}}$:

$$\mathbf{H} = \mathbf{Q} \mathbf{D} \mathbf{Q}^\top, \quad \mathbf{H}^{\frac{1}{2}} = \mathbf{Q} \mathbf{D}^{\frac{1}{2}} \mathbf{Q}^\top,$$

where $\mathbf{Q} = (\mathbf{q}_1 \mathbf{q}_2 \dots \mathbf{q}_d)$ is an orthogonal matrix and $\mathbf{D} = \text{diag}(\lambda_1, \dots, \lambda_d)$, where $\lambda_1 \geq \lambda_2 \geq \dots \lambda_d \geq 0$ are ordered eigenvalues. We assume $\text{rank}(\mathbf{H}) = r$ and $\lambda_r > 0 = \lambda_{r+1}$. Then, from the property of Kronecker product (eg. see [27]), the eigenvalues of $\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}}$ are $\{\sqrt{\lambda_i \lambda_j} : 1 \leq i, j \leq d\}$. The eigenvectors are $\{\mathbf{q}_i \otimes \mathbf{q}_j : 1 \leq i, j \leq d\}$, which forms an orthogonal unit basis of $\mathbb{R}^{d \times d}$. We define

$$\mathcal{S} := \text{span}\{\mathbf{q}_i \otimes \mathbf{q}_j : 1 \leq i, j \leq r\}$$

as the eigenspace spanned by all eigenvectors corresponding to positive eigenvalues. For any matrix $\mathbf{Z}$, we can decompose it as

$$\mathbf{Z} = \sum_{i,j} a_{i,j} \cdot \mathbf{q}_i \otimes \mathbf{q}_j,$$

so

$$(\mathbf{H} \otimes \mathbf{H}) \circ \mathbf{Z} = \sum_{i,j=1}^d a_{i,j} (\mathbf{H} \otimes \mathbf{H}) \circ (\mathbf{q}_i \otimes \mathbf{q}_j) = \sum_{i,j} (\mathbf{H} \mathbf{q}_i) \otimes (\mathbf{H} \mathbf{q}_j)$$

$$= \sum_{i,j=1}^d a_{i,j} \lambda_i \lambda_j \mathbf{q}_i \otimes \mathbf{q}_j = \sum_{i,j=1}^r a_{i,j} \lambda_i \lambda_j \mathbf{q}_i \otimes \mathbf{q}_j \in \mathcal{S},$$

which implies $\mathcal{Z} \subset \mathcal{S}$. On the other hand, for any $\mathbf{Z} \in \mathcal{S}$, we can write it as $\mathbf{Z} = \sum_{i,j=1}^r b_{i,j} \cdot \mathbf{q}_i \otimes \mathbf{q}_j$, so we have

$$\mathbf{Z} = \sum_{i,j=1}^r \frac{b_{i,j}}{\lambda_i \lambda_j} \cdot (\lambda_i \mathbf{q}_i) \otimes (\lambda_j \mathbf{q}_j) = \sum_{i,j=1}^r \frac{b_{i,j}}{\lambda_i \lambda_j} \cdot (\mathbf{H} \mathbf{q}_i) \otimes (\mathbf{H} \mathbf{q}_j) = (\mathbf{H} \otimes \mathbf{H}) \circ \left[\sum_{i,j=1}^r \frac{b_{i,j}}{\lambda_i \lambda_j} (\mathbf{q}_i \otimes \mathbf{q}_j) \right] \in \mathcal{Z},$$

which implies $\mathcal{S} \subset \mathcal{Z}$. Combining two directions, we have $\mathcal{S} = \mathcal{Z}$. Therefore, for any matrix $\mathbf{Z} \in \mathbb{R}^{d \times d}$, we can write it as $\mathbf{Z} = \sum_{i,j=1}^d c_{i,j} \cdot \mathbf{q}_i \otimes \mathbf{q}_j$ for some $c_{i,j}$. Then, we have

$$\begin{aligned} \left(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}} \right) \circ \mathcal{P}_{\mathcal{Z}}(\mathbf{Z}) &= \left(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}} \right) \circ \mathcal{P}_{\mathcal{S}}(\mathbf{Z}) = \left(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}} \right) \circ \sum_{i,j=1}^r c_{i,j} \cdot \mathbf{q}_i \otimes \mathbf{q}_j \\ &= \sum_{i,j=1}^r c_{i,j} \left(\mathbf{H}^{\frac{1}{2}} \mathbf{q}_i \right) \otimes \left(\mathbf{H}^{\frac{1}{2}} \mathbf{q}_j \right) = \sum_{i,j=1}^r c_{i,j} \sqrt{\lambda_i \lambda_j} \cdot \mathbf{q}_i \otimes \mathbf{q}_j. \end{aligned}$$

Therefore, we have

$$\begin{aligned} & \left(\left(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}} \right) \circ \mathcal{P}_{\mathcal{Z}}(\mathbf{Z}) \right) \left(\left(\mathbf{H}^{\frac{1}{2}} \otimes \mathbf{H}^{\frac{1}{2}} \right) \circ \mathcal{P}_{\mathcal{Z}}(\mathbf{Z}) \right)^{\top} \\ &= \left(\sum_{i,j=1}^r c_{i,j} \sqrt{\lambda_i \lambda_j} \cdot \mathbf{q}_i \otimes \mathbf{q}_j \right) \left(\sum_{k,l=1}^r c_{k,l} \sqrt{\lambda_k \lambda_l} \cdot \mathbf{q}_k \otimes \mathbf{q}_l \right)^{\top} \\ &= \sum_{i,j=1}^r c_{i,j}^2 \lambda_i \lambda_j (\mathbf{q}_i \otimes \mathbf{q}_j) (\mathbf{q}_i \otimes \mathbf{q}_j)^{\top} \quad (\text{By orthogonality}) \\ &\succeq \lambda_{-1}^2 \sum_{i,j=1}^r c_{i,j}^2 (\mathbf{q}_i \otimes \mathbf{q}_j) (\mathbf{q}_i \otimes \mathbf{q}_j)^{\top} \\ &= \lambda_{-1}^2 \mathcal{P}_{\mathcal{Z}}(\mathbf{Z}) \cdot \mathcal{P}_{\mathcal{Z}}(\mathbf{Z})^{\top}. \end{aligned}$$

Therefore, we conclude. $\square$

I Experiment Details

Our experiments on GPT2 mostly follow the setting in [14], except that we use the token matrix defined in (3.1). In our experiments, we use $d = 20$, $M = 40$, $\Psi = \mathbf{H} = \mathbf{I}_d$ and $\sigma = 0$. We train the GPT2 with and without MLP layers on two settings: $\beta^* = (0, 0, \dots, 0)^{\top}$ and $\beta^* = (10, 10, \dots, 10)^{\top}$. For GPT2 with and without MLP layers, we initialize the $\mathbf{W}_V$ and $\mathbf{W}_2$ by normal distribution with standard deviation $0.02/\sqrt{\text{number of residual connections}}$, where the number of residual connections equals the number of layers for GPT2 without MLP layers. For GPT2 with MLP layers, this is 2 times the number of layers. Initialization for other matrices follow the default setting in [36]. We use Adam with learning rate 0.0001. We train the model for 200000 steps and we sample a batch of 256 new tasks for each step.

J Does scratchpad help?

In this section, we show the limitations of adding a scratchpad to the token matrix. We will show that by using a single LSA layer and the token matrix with scratchpad, one cannot recover the GD- β estimator defined in Section 5. We leave it as future work to see whether the token matrix with scratchpad could implement other types of estimators that more effectively address the linear regression tasks defined in Assumption 3.1, as well as whether additional structures could help alleviate this approximation error. We follow the notations in previous sections. The token matrix with scratchpad is defined as

$$\mathbf{E} := \begin{pmatrix} \mathbf{X}^\top & \mathbf{x} \\ \mathbf{1}_M^\top & 1 \\ \mathbf{y}^\top & 0 \end{pmatrix} \in \mathbb{R}^{(d+2) \times (M+1)}. \quad (\text{J.1})$$

where $\mathbf{1}_M \in \mathbb{R}^M$ refers to the vector filled with ones.

A LSA model f , is defined by

$$f(\mathbf{E}) := \left[\mathbf{E} + \mathbf{W}_P^\top \mathbf{W}_V \mathbf{E} \mathbf{M} \frac{\mathbf{E}^\top \mathbf{W}_K^\top \mathbf{W}_Q \mathbf{E}}{M} \right]_{-1, -1} = \left[\mathbf{W}_P^\top \mathbf{W}_V \mathbf{E} \mathbf{M} \frac{\mathbf{E}^\top \mathbf{W}_K^\top \mathbf{W}_Q \mathbf{E}}{M} \right]_{-1, -1},$$

where the second equality is because the bottom right entry of $\mathbf{E}$ is zero (see (3.1)). Note that the prediction is the bottom right entry of the output matrix. So only the last row of $\mathbf{W}_P^\top \mathbf{W}_V$ and the last column of $\mathbf{E}$ attend the prediction. Denote

$$\mathbf{W}_P^\top \mathbf{W}_V = \begin{pmatrix} * & * & * \\ \mathbf{w}^\top & \mathbf{a}_1 & \mathbf{a}_2 \end{pmatrix}, \quad \mathbf{W}_K^\top \mathbf{W}_Q = \begin{pmatrix} \mathbf{Q} & \mathbf{b} & * \\ \mathbf{q}_1^\top & \mathbf{b}_1 & * \\ \mathbf{q}_2^\top & \mathbf{b}_2 & * \end{pmatrix},$$

where

$$\mathbf{w} \in \mathbb{R}^d, \mathbf{a}_1, \mathbf{a}_2 \in \mathbb{R}, \mathbf{Q} \in \mathbb{R}^{d \times d}, \mathbf{b}, \mathbf{q}_1, \mathbf{q}_2 \in \mathbb{R}^d, \mathbf{b}_1, \mathbf{b}_2 \in \mathbb{R},$$

and $*$ denotes entries that do not enter the final prediction. Then we have

$$\begin{aligned} f(\mathbf{E}) &= (\mathbf{w}^\top \quad \mathbf{a}_1 \quad \mathbf{a}_2) \frac{\mathbf{E} \mathbf{M} \mathbf{E}^\top}{M} \begin{pmatrix} \mathbf{Q} & \mathbf{b} & * \\ \mathbf{q}_1^\top & \mathbf{b}_1 & * \\ \mathbf{q}_2^\top & \mathbf{b}_2 & * \end{pmatrix} \begin{pmatrix} \mathbf{x} \\ 1 \\ 0 \end{pmatrix} \\ &= (\mathbf{w}^\top \quad \mathbf{a}_1 \quad \mathbf{a}_2) \frac{\mathbf{E} \mathbf{M} \mathbf{E}^\top}{M} \begin{pmatrix} \mathbf{Q} & \mathbf{b} \\ \mathbf{q}_1^\top & \mathbf{b}_1 \\ \mathbf{q}_2^\top & \mathbf{b}_2 \end{pmatrix} \begin{pmatrix} \mathbf{x} \\ 1 \end{pmatrix} \\ &= (\mathbf{w}^\top \quad \mathbf{a}_1 \quad \mathbf{a}_2) \frac{1}{M} \begin{pmatrix} \mathbf{X}^\top \mathbf{X} & \mathbf{X}^\top \mathbf{1}_M & \mathbf{X}^\top \mathbf{y} \\ \mathbf{1}_M^\top \mathbf{X} & M & \mathbf{1}_M^\top \mathbf{y} \\ \mathbf{y}^\top \mathbf{X} & \mathbf{y}^\top \mathbf{1}_M & \mathbf{y}^\top \mathbf{y} \end{pmatrix} \begin{pmatrix} \mathbf{Q} & \mathbf{b} \\ \mathbf{q}_1^\top & \mathbf{b}_1 \\ \mathbf{q}_2^\top & \mathbf{b}_2 \end{pmatrix} \begin{pmatrix} \mathbf{x} \\ 1 \end{pmatrix}. \end{aligned}$$

Following the Assumption 3.1, we use $\tilde{\beta}$ to refer to the task parameter, then

$$\tilde{\beta} := \beta^* + \Psi^{\frac{1}{2}} \tilde{\theta}, \quad \text{where } \tilde{\theta} \sim \mathcal{N}(\mathbf{0}_d, \mathbf{I}_d).$$

We then decompose $f(\mathbf{E})$ as

$$f(\mathbf{E}) = \underbrace{(\mathbf{w}^\top \quad \mathbf{a}_1 \quad \mathbf{a}_2) \frac{1}{M} \begin{pmatrix} \mathbf{X}^\top \mathbf{X} & \mathbf{X}^\top \mathbf{1}_M & \mathbf{X}^\top \mathbf{y} \\ \mathbf{1}_M^\top \mathbf{X} & M & \mathbf{1}_M^\top \mathbf{y} \\ \mathbf{y}^\top \mathbf{X} & \mathbf{y}^\top \mathbf{1}_M & \mathbf{y}^\top \mathbf{y} \end{pmatrix} \begin{pmatrix} \mathbf{Q} \\ \mathbf{q}_1^\top \\ \mathbf{q}_2^\top \end{pmatrix}}_{\text{I}} \cdot \mathbf{x} + \text{II},$$

where terms I and II are independent of $\mathbf{x}$. Therefore, we have

$$\begin{aligned} \mathcal{R}(f) &:= \mathbb{E} (f(\mathbf{E}) - y)^2 \\ &= \mathbb{E} \left(f(\mathbf{E}) - \langle \tilde{\beta}, \mathbf{x} \rangle \right)^2 + \sigma^2 && \text{since } y | \tilde{\beta}, \mathbf{x} \sim \mathcal{N}(0, \sigma^2) \\ &= \mathbb{E} \left(\text{I} \cdot \mathbf{x} + \text{II} - \langle \tilde{\beta}, \mathbf{x} \rangle \right)^2 + \sigma^2 && \text{since } f(\mathbf{E}) = \text{I} \cdot \mathbf{x} + \text{II} \\ &= \mathbb{E} \|\text{I}^\top - \tilde{\beta}\|_{\mathbf{H}}^2 + \mathbb{E} \text{II}^2 + \sigma^2 && \text{since } \mathbf{x} \sim \mathcal{N}(0, \mathbf{H}) \end{aligned}$$

$$\geq \mathbb{E} \|\mathbf{I}^\top - \tilde{\boldsymbol{\beta}}\|_{\mathbf{H}}^2 + \sigma^2.$$

Note that the above equation holds if $\mathbf{\Pi} = 0$, which can be easily achieved by setting $\mathbf{b} = \mathbf{0}_d, \mathbf{b}_1 = \mathbf{b}_2 = 0$. Therefore, without loss of generality, we can consider the LSA function which takes the following form:

$$f(\mathbf{E}) = \mathbf{I} \cdot \mathbf{x} = \begin{pmatrix} \mathbf{w}^\top & \mathbf{a}_1 & \mathbf{a}_2 \end{pmatrix} \frac{1}{M} \begin{pmatrix} \mathbf{X}^\top \mathbf{X} & \mathbf{X}^\top \mathbf{1}_M & \mathbf{X}^\top \mathbf{y} \\ \mathbf{1}_M^\top \mathbf{X} & M & \mathbf{1}_M^\top \mathbf{y} \\ \mathbf{y}^\top \mathbf{X} & \mathbf{y}^\top \mathbf{1}_M & \mathbf{y}^\top \mathbf{y} \end{pmatrix} \begin{pmatrix} \mathbf{Q} \\ \mathbf{q}_1^\top \\ \mathbf{q}_2^\top \end{pmatrix} \cdot \mathbf{x}. \quad (\text{J.2})$$

Let's then try to determine whether such a function can effectively represent a GD- $\boldsymbol{\beta}$ function, which takes the form of

$$f_{\text{GD-}\boldsymbol{\beta}}(\mathbf{E}) = \langle \boldsymbol{\beta}^*, \mathbf{x} \rangle - \left\langle \frac{\boldsymbol{\Gamma}^* \mathbf{X}^\top (\mathbf{X} \boldsymbol{\beta}^* - \mathbf{y})}{M}, \mathbf{x} \right\rangle, \quad (\text{J.3})$$

where $\boldsymbol{\beta}^*$ is the prior mean in Assumption 3.1 in our submission, and $\boldsymbol{\Gamma}^* = \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Omega}^{-1} \mathbf{H}^{\frac{1}{2}}$ and $\boldsymbol{\Omega} = \frac{M+1}{M} \mathbf{H}^{\frac{1}{2}} \boldsymbol{\Psi} \mathbf{H}^{\frac{1}{2}} + \frac{\sigma^2 + \text{tr}(\mathbf{H} \boldsymbol{\Psi})}{M} \mathbf{I}_d$ is defined in (5.3) in our submission. In order to achieve this, one simple way is to let

$$\mathbf{w} = -\boldsymbol{\beta}^*, \quad \mathbf{a}_1 = \mathbf{a}_2 = \mathbf{1}, \quad \mathbf{Q} = (\boldsymbol{\Gamma}^*)^\top, \quad \mathbf{q}_1 = \boldsymbol{\beta}^*, \quad \mathbf{q}_2 = \mathbf{0}_d. \quad (\text{J.4})$$

However, this will incur some additive terms and will potentially enlarge the ICL risk. Inserting the parameters above, we have

$$f(\mathbf{E}) = f_{\text{GD-}\boldsymbol{\beta}}(\mathbf{E}) + \frac{\mathbf{1}_M^\top}{M} (\mathbf{X}(\boldsymbol{\Gamma}^*)^\top - \mathbf{x} \boldsymbol{\beta}^* (\boldsymbol{\beta}^*)^\top + \mathbf{y} (\boldsymbol{\beta}^*)^\top) \cdot \mathbf{x}$$

This shows that the extended token with a single LSA cannot easily implement the GD- $\boldsymbol{\beta}$ function class and may incur an additive ICL risk depending on $\boldsymbol{\beta}^*$ (as shown in Theorem 4.1).

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: We believe our introduction and abstract are factually accurate in describing the contributions of the paper and of its result.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We discuss the limitation in Section 4 and the future work section.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory Assumptions and Proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: We provide the assumption in Assumption 3.1.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental Result Reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We provide all detailed experiment setup in Appendix I.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: The data we use are simulated. We plan to release our code upon acceptance.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental Setting/Details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: We include all details and hyperparameters in our experiments in Appendix I.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment Statistical Significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [No]

Justification: The cost for producing the experiments precludes us from reporting error bars. Moreover, the experiment section is not the main contribution of this paper.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).

- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments Compute Resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [NA]

Justification: There is no strict requirement for the computer resources for our experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code Of Ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: None to report.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader Impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: This work focuses on the theory of in-context learning so it is not expected to have a direct societal impact.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.

- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The data used in this paper are fully simulated, so there is no data or model that have a high risk.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [NA]

Justification: We do not release any asset.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.

- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New Assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: We do not release any asset.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and Research with Human Subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: We do not involve humans in our research.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional Review Board (IRB) Approvals or Equivalent for Research with Human Subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: Our paper does not involve crowdsourcing.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.