
High-dimensional Analysis of Synthetic Data Selection

Parham Rezaei* Filip Kovacevic* Francesco Locatello*† Marco Mondelli*†
Institute of Science and Technology (ISTA)

Abstract

Despite the progress in the development of generative models, their usefulness in creating synthetic data that improve prediction performance of classifiers has been put into question. Besides heuristic principles such as “synthetic data should be close to the real data distribution”, it is actually not clear which specific properties affect the generalization error. Our paper addresses this question through the lens of high-dimensional regression. Theoretically, we show that, for linear models, the *covariance shift* between the target distribution and the distribution of the synthetic data affects the generalization error but, surprisingly, the mean shift does not. Furthermore we prove that, in some settings, matching the covariance of the target distribution is optimal. Remarkably, the theoretical insights from linear models carry over to deep neural networks and generative models. We empirically demonstrate that the *covariance matching* procedure (matching the covariance of the synthetic data with that of the data coming from the target distribution) performs well against several recent approaches for synthetic data selection, across training paradigms, architectures, datasets and generative models used for augmentation.

1 Introduction

The controllable generation of arbitrary amounts of synthetic data for training machine learning models has long been considered as one of the key implications unlocked by more capable generative models [Kingma and Welling, 2014, Goodfellow et al., 2014, Shrivastava et al., 2017, Nikolenko et al., 2021]. After all, synthetic data can not only be abundant, which would already be tremendously impactful in data-scarce applications such as medicine [Esteban et al., 2017, van Breugel et al., 2024], but it can also address other difficulties of observational data, such as privacy [Jordon et al., 2018], imbalancedness [Parihar et al., 2024, Ramaswamy et al., 2021] and overall difficulty to collect, as the domain can be specific [Dunlap et al., 2023] or the task complex [Wang et al., 2023]. At the same time, while generative models have progressed significantly, experimental results are still mixed. Prior studies highlight both the promise [Trabucco et al., 2024, He et al., 2023, Azizi et al., 2023, Dunlap et al., 2023] and pitfalls [Fan et al., 2024, Burg et al., 2023, Geng et al., 2024, Shumailov et al., 2024, Wyllie et al., 2024] of synthetic data. What emerges here is a broad challenge which consists of understanding *how extra synthetic data, for example from a generative model, helps training predictors*. Our paper tackles this challenge theoretically and empirically. To do so, we assume access to a training dataset (X_t, y_t) that contains i.i.d. samples, as well as to an additional *synthetic dataset* (X_s, y_s) . The samples from the synthetic dataset are also i.i.d., but they come from a different distribution, since they are obtained from a generative model and not from the training dataset. We perform empirical risk minimization (ERM) using the augmentation $((X_t, X_s), (y_t, y_s))$, and evaluate the performance on an independent test sample with the same distribution as (X_t, y_t) . In this context, the challenge above leads to the following concrete question:

How to select the dataset (X_s, y_s) in order to minimize the test error? (Q)

*Emails: parhamix@gmail.com, {filip.kovacevic, francesco.locatello, marco.mondelli}@ist.ac.at

†Equal advising

By studying this question, we can identify which properties of the distribution of (X_s, y_s) improve generalization, thus guiding the selection of data obtained in practice from generative models.

Formalization of the problem. We assume that the distributions of both the original training dataset and the additional synthetic one are mixture models. The number of mixtures corresponds to the number of classes in the datasets, with each mixture component corresponding to a single class. As common in practice [Burg et al., 2023], the data augmentation occurs class-by-class. We then address the question (Q) when (X_t, y_t) and (X_s, y_s) correspond to a single class. While this is a strong assumption chosen for mathematical tractability, we highlight that the resulting data selection procedure is extensively tested in practical settings where it performs well against existing baselines. For related work, see Appendix A.

Main contributions. Our analysis reveals a *surprising* behavior: while the shift in the covariances of X_t and X_s affects the test error, the *shift in the means of X_t and X_s is irrelevant*, as long as the training dataset (X_t, y_t) is not too small compared to the synthetic dataset (X_s, y_s) . This allows synthetic data selection to be formulated as an optimization over the covariance Σ_s of X_s , and *covariance matching* ($\Sigma_s \propto \Sigma_t$, where Σ_t is the covariance of X_t) yields optimal performance in some settings. Empirically, relying on this principle alone performs on par—or even outperforms—several recent approaches for synthetic data selection. We summarize our contributions below:

- We precisely characterize the test error of the min-norm least squares estimator when the dimensions of β , y_t , and y_s grow proportionally. In both under- and over-parameterized regimes (Theorems 3.1 and D.1), the test error converges to a deterministic limit depending *only on the covariances* Σ_t, Σ_s , and not on the means μ_t, μ_s .
- Our characterization implies that we can select synthetic data minimizing the test error based on their covariance. We then show that, under some conditions, taking $\Sigma_s \propto \Sigma_t$, i.e., *covariance matching*, is optimal (Theorems 3.2 and D.2 for under-parameterized and over-parameterized regimes).
- We validate the effectiveness of covariance matching as a way to select synthetic data in several practical scenarios. This simple approach matches or surpasses the recent baselines [He et al., 2023, Lin et al., 2023, Hulkund et al., 2025] across various training paradigms (training from scratch, distilling a bigger model, fine-tuning a model trained from a larger dataset), architectures (ResNets, transformers), datasets (CIFAR-10, ImageNet-100, RxRx1), and generative models used to obtain synthetic data (StyleGAN2-Ada, SANA1.5, PixArt- α , StableDiffusion1.4, MorphGen).

2 Preliminaries

Data model. We consider data augmentation in the context of linear models. Formally, we observe two datasets (X_t, y_t) and (X_s, y_s) , denoting training data and augmenting synthetic data, such that

$$y_{(i)} = X_{(i)}\beta + \varepsilon_{(i)}, \quad (i) \in \{t, s\}, \quad (2.1)$$

where $X_{(i)} \in \mathbb{R}^{n_{(i)} \times p}$, $\beta \in \mathbb{R}^p$, and $\varepsilon_{(i)} \in \mathbb{R}^{n_{(i)}}$. Thus, we are given n_t training samples and n_s synthetic samples, all of which are p dimensional. We denote the total number of samples as $n := n_t + n_s$. For $(i) \in \{t, s\}$, each entry of the noise vector $\varepsilon_{(i)}$ is i.i.d. with mean zero and variance σ^2 , and the rows of $X_{(i)}$ are independent random vectors with $p \times p$ population covariance matrix $\Sigma_{(i)}$ and mean $\mu_{(i)}$. By omitting subscripts, we introduce $X := [X_t \ X_s]^\top \in \mathbb{R}^{n \times p}$, $y := [y_t \ y_s]^\top \in \mathbb{R}^n$. We will consider our results under mild assumptions on X (see Appendix B).

Risk and estimator. We test estimators on data sampled from the same distribution as the training dataset (X_t, y_t) and, given an estimator $\hat{\beta}$, its out-of-sample excess risk is defined as

$$R_X(\hat{\beta}; \beta) := \mathbb{E}[(x_t^\top \hat{\beta} - x_t^\top \beta)^2 \mid X] = \mathbb{E}[\|\hat{\beta} - \beta\|_{\Sigma_t + \mu_t \mu_t^\top}^2 \mid X], \quad (2.2)$$

where x_t has the same distribution as a row of X_t and $\|x\|_M^2 := x^\top M x$. Specifically, we are interested in the performance of the min-norm least squares regression estimator of y on the whole dataset available X , i.e., $\hat{\beta} := \arg \min \{\|b\|_2 : b \text{ minimizes } \|y - Xb\|_2^2\}$. This estimator is also motivated by its close relation to the gradient descent optimum.³

³Gradient descent converges to the interpolator closest in ℓ_2 norm to the initialization (see Equation (33) in Bartlett et al. [2021]) and, as such, $\hat{\beta}$ corresponds to the gradient descent solution starting from 0 initialization.

3 Theoretical results

We characterize the excess risk of the min-norm interpolator using both training and augmenting synthetic data. We then use the explicitly derived formulas to optimize the data selection process, in which, surprisingly, distribution means play no role. We contrast this setting with having only synthetic data available, where means instead impact the excess risk. Our findings hold in both the under-parameterized and over-parameterized regimes.

For clarity, we present the under-parameterized regime here, and defer the analysis of the over-parameterized regime to Appendix D. Accordingly, we assume that $n/p, n_t/p, n_s/p \in [1 + \tau, 1/\tau]$, for some small $\tau > 0$. This implies that $n > p$, indeed making the setting under-parameterized. The following result provides a precise asymptotic characterization of the excess risk, extending the results by Yang et al. [2025] to non-zero centered data. Its proof is deferred to Appendix E.2.

Theorem 3.1. *Let $M = \Sigma_t^{-1/2} \Sigma_s \Sigma_t^{-1/2}$ and denote its eigenvalues as $\lambda_1 \geq \dots \geq \lambda_p$. Then, under the assumptions from Section 2 and the start of this section, it holds that, with high probability,*

$$\lim_{n \rightarrow \infty} \left| R_X(\hat{\beta}; \beta) - \frac{\sigma^2}{n} \text{Tr} \left[(\alpha_1 M + \alpha_2 I_p)^{-1} \right] \right| = 0, \quad (3.1)$$

where α_1 and α_2 are the unique positive solutions to the following two equations

$$\alpha_1 + \alpha_2 = 1 - \frac{p}{n}, \quad \alpha_1 + \frac{1}{n} \sum_{i=1}^p \frac{\lambda_i \alpha_1}{\lambda_i \alpha_1 + \alpha_2} = \frac{n_s}{n}. \quad (3.2)$$

Theorem 3.1 gives a *deterministic equivalent* of the test error obtained using training and synthetic data in the under-parameterized regime. In fact, $R_X(\hat{\beta}; \beta)$ is a random quantity (the data is random), while $\frac{\sigma^2}{n} \text{Tr}[(\alpha_1 M + \alpha_2 I_p)^{-1}]$ is deterministic as it depends on properties of the data distributions. Remarkably, the deterministic equivalent depends only on the covariances Σ_t, Σ_s (via $M = \Sigma_t^{-1/2} \Sigma_s \Sigma_t^{-1/2}$) and it does not depend on the means μ_t, μ_s . The independence of the test error on the mean shift is surprising, and it is in stark contrast with the setting in which we only train on (X_s, y_s) , where the performance does depend on μ_s, μ_t (see Appendix C.1).

The deterministic equivalent can be optimized to find the covariance Σ_s minimizing the error. Towards that end, let us denote the deterministic quantity from (3.1) as

$$\mathcal{R}_u(M) := \frac{\sigma^2}{n} \text{Tr} \left[(\alpha_1 M + \alpha_2 I_p)^{-1} \right], \quad (3.3)$$

where α_1 and α_2 satisfy (3.2). This corresponds to the limit of the risk $R_X(\hat{\beta}; \beta)$ due to Theorem 3.1. Thus, the guiding question (Q) posed in the introduction can be formalized as:

Given Σ_t , what is the optimal Σ_s that minimizes $\mathcal{R}_u(\Sigma_t^{-1/2} \Sigma_s \Sigma_t^{-1/2})$?

The following theorem exactly treats this. Its proof is in Appendix E.4.

Theorem 3.2. *Let $\mathcal{M} := \{M \in \mathbb{R}_{>0}^{p \times p} : \text{Tr}[M] = p\}$, where $\mathbb{R}_{>0}^{p \times p}$ denotes the set of $p \times p$ positive definite matrices. Then, it holds that $I_p = \arg\inf_{M \in \mathcal{M}} \mathcal{R}_u(M)$.*

Theorem 3.2 proves that, having fixed $\text{Tr}[M]$, the limit risk $\mathcal{R}_u(M)$ is minimized for M proportional to I_p . The imposed trace normalization is justified in Appendix C.2. Thus, given a training covariance Σ_t , choosing synthetic data with $\Sigma_s \propto \Sigma_t$, i.e., *matching the covariances*, is optimal.

4 Experimental results

Our theoretical results in Section 3 and Appendix D show the optimality of *covariance matching* ($\Sigma_s \propto \Sigma_t$) in both under-parameterized and over-parameterized regimes. We now consider classification problems, assume access to a large pool of synthetic samples obtained from generative models, and perform the augmentation per class. We implement *covariance matching* via a greedy algorithm: we initialize $S = \emptyset$ and, until $|S| = n_s$, we add the x from the generated pool that minimizes $\|\hat{\Sigma}(S \cup \{x\}) - \hat{\Sigma}_t\|_F$, where $\hat{\Sigma}(\cdot)$ and $\hat{\Sigma}_t$ denote the sample covariance of CLIP features of the synthetic samples and real samples respectively and $\|\cdot\|_F$ is the Frobenius norm. To accelerate the selection, we compute covariances in a 32-dimensional PCA space fit on the n_t real reference features. After the selection, we train a classifier on the union of real and selected synthetic samples.

Experimental setup. When using CIFAR-10, we evaluate three training paradigms. (1) *Scratch*: train a ResNet-18 [He et al., 2016] from scratch on the available data. (2) *Distillation*: train a ResNet-

18 using soft targets (logits) from a ResNet-50 trained on full CIFAR-10, following Hinton et al. [2015]. (3) *Pretrained*: fine-tune an ImageNet-pretrained ResNet-18 with a new classification head. We also repeat the *Scratch* and *Distillation* experiments replacing the ResNet with two transformer models (ViT and Swin-T). Unless stated otherwise, we use $n_t = 200$ real images and augment with $n_s = 800$ synthetic images per class. The features for the selection algorithms are extracted with CLIP ViT-B, yielding a $p = 512$ -dimensional feature space, which places us in an under-parameterized regime. We report in Table 8 in Appendix F an additional experiment in the over-parameterized regime. We additionally consider ImageNet-100 as a more diverse dataset and RxRx1 [Sypetkowski et al., 2023] as a specialized one, see again Appendix F for details.

Baselines. We compare *Covariance matching* with the following baselines. (1) *Center matching* [He et al., 2023], (2) *Center sampling* [Lin et al., 2023], (3) *DS3* [Hulkund et al., 2025], (4) *K-means* [Lin et al., 2023], (5) *Random* (equivalent to methods “No-filtering” [Hulkund et al., 2025], “Match-dist” [Hulkund et al., 2025], and “Match-label” [Hulkund et al., 2025] due to having the same number of data for each class), (6) *Text matching* [Lin et al., 2023], and (7) *Text sampling* [Lin et al., 2023]. We also include *No synthetic* (using only n_t real samples) and *Real upper bound* (using $n_t + n_s$ real samples). Results are averaged over 10 seeds (5 for Table 2a in Appendix F) and reported as mean ± 1 standard deviation. Details of all baselines are provided in Appendix F.

Main findings. First, we test *diversity/quality trade-offs*. To do so, for each class we generate images with StyleGAN2-Ada [Karras et al., 2020] under different truncations [Karras et al., 2019]: 6K images from a 0.2-truncated model with three randomized truncation centers and 4K images from a 0.6-truncated model with two randomized centers. This produces synthetic data with varying diversity and fidelity. The results of Table 1 demonstrate that covariance matching outperforms all baselines for all training paradigms. Table 9 in Appendix F suggests that this superiority is partly due to selecting more diverse samples, evident from the improved Recall [Kynkäänniemi et al., 2019], FID [Heusel et al., 2017], and KID [Bińkowski et al., 2018] scores guaranteed by covariance matching. We also demonstrate the effectiveness for transformer models in Table 4 in Appendix F.

Second, we test *text-to-image (T2I) generative models*. To do so, for each class we generate 4K SANA-1.5 [Xie et al., 2025], 4K PixArt- α [Chen et al., 2024], and 2K StableDiffusion1.4 [Rombach et al., 2022] images. Table 1 shows that covariance matching also performs well in this mixed setup. Finally, to demonstrate the generality of our findings, we consider a broader dataset from computer vision (ImageNet-100) and a specialized dataset from fluorescence microscopy (RxRx1, [Sypetkowski et al., 2023]). Once again, the results reported in Tables 2a-2b in Appendix F show that covariance matching performs on par with the best baselines in all settings.

Table 1: *Covariance matching* outperforms all baselines across three training paradigms on CIFAR-10, when the synthetic data is generated via five truncated StyleGAN2-Ada and various T2I models.

Method	Truncated generators			T2I generators		
	Scratch	Distillation	Pretrained	Scratch	Distillation	Pretrained
No synthetic	44.36 $\pm$ 1.51	47.33 $\pm$ 0.57	63.40 $\pm$ 1.33	44.36 $\pm$ 1.51	47.33 $\pm$ 0.57	63.40 $\pm$ 1.33
Center matching [He et al., 2023]	50.04 $\pm$ 2.84	53.83 $\pm$ 0.59	67.01 $\pm$ 0.89	53.46 $\pm$ 1.95	57.67 $\pm$ 0.58	66.52 $\pm$ 0.81
Center sampling [Lin et al., 2023]	50.48 $\pm$ 2.03	54.91 $\pm$ 1.07	67.71 $\pm$ 0.90	50.15 $\pm$ 1.79	56.05 $\pm$ 0.65	65.38 $\pm$ 0.98
DS3 [Hulkund et al., 2025]	52.83 $\pm$ 2.19	58.32 $\pm$ 0.43	68.21 $\pm$ 0.66	54.15 $\pm$ 2.17	59.43 $\pm$ 0.73	66.00 $\pm$ 0.94
K-means [Lin et al., 2023]	50.74 $\pm$ 1.77	56.06 $\pm$ 0.68	66.50 $\pm$ 1.11	51.63 $\pm$ 1.29	56.77 $\pm$ 0.89	65.23 $\pm$ 0.61
Random	49.38 $\pm$ 2.43	54.89 $\pm$ 0.91	67.65 $\pm$ 0.77	51.26 $\pm$ 1.96	55.27 $\pm$ 0.74	65.24 $\pm$ 1.01
Text matching [Lin et al., 2023]	50.94 $\pm$ 1.40	55.17 $\pm$ 0.57	67.81 $\pm$ 0.76	51.20 $\pm$ 1.82	56.08 $\pm$ 0.57	65.93 $\pm$ 0.59
Text sampling [Lin et al., 2023]	50.28 $\pm$ 1.18	54.82 $\pm$ 0.72	67.45 $\pm$ 1.02	50.31 $\pm$ 1.70	55.79 $\pm$ 0.68	64.93 $\pm$ 1.12
Covariance matching (ours)	54.00 $\pm$ 1.89	59.77 $\pm$ 0.61	69.20 $\pm$ 0.56	54.45 $\pm$ 2.11	59.17 $\pm$ 0.64	66.69 $\pm$ 0.70
Real upper bound	61.08 $\pm$ 2.54	65.38 $\pm$ 0.51	74.35 $\pm$ 0.56	61.08 $\pm$ 2.54	65.38 $\pm$ 0.51	74.35 $\pm$ 0.56

5 Conclusion

This paper offers the first step in understanding the precise connection between training on a mix of real and synthetic data and generalizing on real data. We start with a high-dimensional linear regression analysis, where we find that only covariance shifts, and not mean shifts, affect the error. Even if our theory ignores the interactions between classes that would affect neural network training, the resulting insights transfer to realistic settings. We empirically demonstrate that matching the covariance between samples from real image classification datasets and generative models (irrespective of whether they are from GANs or diffusion model variants) improves the accuracy of deep networks (ResNets and Transformers) under different training regimes (from scratch, distillation, and fine-tuning). In fact, our principled approach even performs on-par or better than existing

baselines [Hulkund et al., 2025, He et al., 2023, Lin et al., 2023]. Our analysis provides a foundation for studying real–synthetic data interactions, and extending these results to more complex models and data settings opens several directions for future work (see Appendix G).

References

- S. Azizi, S. Kornblith, C. Saharia, M. Norouzi, and D. J. Fleet. Synthetic data from diffusion models improves imagenet classification. *Transactions on Machine Learning Research*, 2023.
- Z. Bai and J. W. Silverstein. *Spectral analysis of large dimensional random matrices*. Springer, 2010.
- S. Barocas, M. Hardt, and A. Narayanan. Fairness and machine learning. *Recommender systems handbook*, 1:453–459, 2020.
- P. L. Bartlett, P. M. Long, G. Lugosi, and A. Tsigler. Benign overfitting in linear regression. *Proceedings of the National Academy of Sciences*, 117(48):30063–30070, 2020.
- P. L. Bartlett, A. Montanari, and A. Rakhlin. Deep learning: a statistical viewpoint. *Acta numerica*, 30:87–201, 2021.
- M. Belkin, D. Hsu, S. Ma, and S. Mandal. Reconciling modern machine-learning practice and the classical bias–variance trade-off. *Proceedings of the National Academy of Sciences*, 116(32):15849–15854, 2019.
- M. Bińkowski, D. J. Sutherland, M. Arbel, and A. Gretton. Demystifying mmd gans. In *International Conference on Learning Representations*, 2018.
- S. Bombari and M. Mondelli. Spurious correlations in high dimensional regression: The roles of regularization, simplicity bias and over-parameterization. In *International Conference on Machine Learning*, 2025.
- M. F. Burg, F. Wenzel, D. Zietlow, M. Horn, O. Makansi, F. Locatello, and C. Russell. Image retrieval outperforms diffusion models on data augmentation. *Transactions on Machine Learning Research*, 2023.
- R. Cadei, I. Demirel, P. De Bartolomeis, L. Lindorfer, S. Cremer, C. Schmid, and F. Locatello. Causal lifting of neural representations: Zero-shot generalization for causal inferences. *arXiv preprint arXiv:2502.06343*, 2025.
- X. Chang, Y. Li, S. Oymak, and C. Thrampoulidis. Provable benefits of overparameterization in model compression: From double descent to pruning neural networks. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pages 6974–6983, 2021.
- J. Chen, Y. Jincheng, G. Chongjian, L. Yao, E. Xie, Z. Wang, J. Kwok, P. Luo, H. Lu, and Z. Li. Pixart- α : Fast training of diffusion transformer for photorealistic text-to-image synthesis. In *International Conference on Learning Representations*, 2024.
- C. Cheng and A. Montanari. Dimension free ridge regression. *The Annals of Statistics*, 52(6):2879 – 2912, 2024.
- B. Demirel, M. Fumero, T. Karaletsos, and F. Locatello. Morphgen: Controllable and morphologically plausible generative cell-imaging. In *ICML 2025 Workshop on Scaling Up Intervention Models*, 2025.
- L. Dunlap, A. Umno, H. Zhang, J. Yang, J. E. Gonzalez, and T. Darrell. Diversify your vision datasets with automatic diffusion-based augmentation. *Advances in Neural Information Processing Systems*, 36, 2023.
- C. Dwork. Differential privacy. In *International colloquium on automata, languages, and programming*, pages 1–12. Springer, 2006.
- C. Esteban, S. L. Hyland, and G. Rätsch. Real-valued (medical) time series generation with recurrent conditional gans. *arXiv preprint arXiv:1706.02633*, 2017.

- L. Fan, K. Chen, D. Krishnan, D. Katabi, P. Isola, and Y. Tian. Scaling laws of synthetic images for model training... for now. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 7382–7392, 2024.
- S. Y. Gadre, G. Ilharco, A. Fang, J. Hayase, G. Smyrnis, T. Nguyen, R. Marten, M. Wortsman, D. Ghosh, J. Zhang, et al. Datacomp: In search of the next generation of multimodal datasets. *Advances in Neural Information Processing Systems*, 36:27092–27112, 2023.
- S. Geng, C.-Y. Hsieh, V. Ramanujan, M. Wallingford, C.-L. Li, P. W. W. Koh, and R. Krishna. The unmet promise of synthetic training images: Using retrieved real images performs better. *Advances in Neural Information Processing Systems*, 37:7902–7929, 2024.
- I. J. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio. Generative adversarial nets. *Advances in Neural Information Processing Systems*, 27, 2014.
- Q. Han and X. Xu. The distribution of ridgeless least squares interpolators. *arXiv preprint arXiv:2307.02044*, 2023.
- T. Hastie, A. Montanari, S. Rosset, and R. J. Tibshirani. Surprises in high-dimensional ridgeless least squares interpolation. *Annals of statistics*, 50(2):949, 2022.
- K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 770–778, 2016.
- R. He, S. Sun, X. Yu, C. Xue, W. Zhang, P. Torr, S. Bai, and X. Qi. Is synthetic data from generative models ready for image recognition? In *International Conference on Learning Representations*, 2023.
- M. Heusel, H. Ramsauer, T. Unterthiner, B. Nessler, and S. Hochreiter. Gans trained by a two time-scale update rule converge to a local nash equilibrium. *Advances in Neural Information Processing Systems*, 30, 2017.
- G. Hinton, O. Vinyals, and J. Dean. Distilling the knowledge in a neural network. *arXiv preprint arXiv:1503.02531*, 2015.
- R. A. Horn and C. R. Johnson. *Matrix analysis*. Cambridge university press, 2012.
- N. Hulkund, A. Maalouf, L. Cai, D. Yang, T.-H. Wang, A. O’Neil, T. Haucke, S. Mukherjee, V. Ramaswamy, J. H. Shen, et al. Datas³: Dataset subset selection for specialization. *arXiv preprint arXiv:2504.16277*, 2025.
- M. E. Ildiz, H. A. Gozeten, E. O. Taga, M. Mondelli, and S. Oymak. High-dimensional analysis of knowledge distillation: Weak-to-strong generalization and scaling laws. In *International Conference on Learning Representations*, 2025.
- A. Jain, A. Montanari, and E. Sasoglu. Scaling laws for learning with real and surrogate data. *Advances in Neural Information Processing Systems*, 37:110246–110289, 2024.
- J. Jordon, J. Yoon, and M. Van Der Schaar. Pate-gan: Generating synthetic data with differential privacy guarantees. In *International Conference on Learning Representations*, 2018.
- T. Karras, S. Laine, and T. Aila. A style-based generator architecture for generative adversarial networks. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 4401–4410, 2019.
- T. Karras, M. Aittala, J. Hellsten, S. Laine, J. Lehtinen, and T. Aila. Training generative adversarial networks with limited data. *Advances in Neural Information Processing Systems*, 33:12104–12114, 2020.
- D. P. Kingma and M. Welling. Auto-encoding variational bayes. In *International Conference on Learning Representations*, 2014.
- G. Kolossov, A. Montanari, and P. Tandon. Towards a statistical theory of data selection under weak supervision. In *International Conference on Learning Representations*, 2024.

- T. Kynkäänniemi, T. Karras, S. Laine, J. Lehtinen, and T. Aila. Improved precision and recall metric for assessing generative models. *Advances in Neural Information Processing Systems*, 32, 2019.
- S. Lin, K. Wang, X. Zeng, and R. Zhao. Explore the power of synthetic data on few-shot object detection. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 638–647, 2023.
- N. R. Mallinar, A. Zane, S. Frei, and B. Yu. Minimum-norm interpolation under covariate shift. In *International Conference on Machine Learning*, 2024.
- A. W. Marshall, I. Olkin, and B. C. Arnold. Inequalities: theory of majorization and its applications. *Springer*, 1979.
- A. Montanari, F. Ruan, Y. Sohn, and J. Yan. The generalization error of max-margin linear classifiers: Benign overfitting and high dimensional asymptotics in the overparametrized regime. *arXiv preprint arXiv:1911.01544*, 2019.
- M. F. Naeem, S. J. Oh, Y. Uh, Y. Choi, and J. Yoo. Reliable fidelity and diversity metrics for generative models. In *International conference on machine learning*, pages 7176–7185. PMLR, 2020.
- S. I. Nikolenko et al. *Synthetic data for deep learning*, volume 174. Springer, 2021.
- J. Nixon, M. W. Dusenberry, L. Zhang, G. Jerfel, and D. Tran. Measuring calibration in deep learning. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops*, June 2019.
- R. Parihar, A. Bhat, A. Basu, S. Mallick, J. N. Kundu, and R. V. Babu. Balancing act: distribution-guided debiasing in diffusion models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 6668–6678, 2024.
- P. Patil, J.-H. Du, and R. J. Tibshirani. Optimal ridge regularization for out-of-distribution prediction. In *International Conference on Machine Learning*, pages 39908–39954, 2024.
- A. Radford, J. W. Kim, C. Hallacy, A. Ramesh, G. Goh, S. Agarwal, G. Sastry, A. Askell, P. Mishkin, J. Clark, et al. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, pages 8748–8763. PMLR, 2021.
- V. V. Ramaswamy, S. S. Kim, and O. Russakovsky. Fair attribute classification through latent space de-biasing. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 9301–9310, 2021.
- D. Richards, J. Mourtada, and L. Rosasco. Asymptotics of ridge (less) regression under general source condition. In *International Conference on Artificial Intelligence and Statistics*, pages 3889–3897. PMLR, 2021.
- R. Rombach, A. Blattmann, D. Lorenz, P. Esser, and B. Ommer. High-resolution image synthesis with latent diffusion models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 10684–10695, 2022.
- A. Shrivastava, T. Pfister, O. Tuzel, J. Susskind, W. Wang, and R. Webb. Learning from simulated and unsupervised images through adversarial training. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 2107–2116, 2017.
- I. Shumailov, Z. Shumaylov, Y. Zhao, N. Papernot, R. Anderson, and Y. Gal. Ai models collapse when trained on recursively generated data. *Nature*, 631(8022):755–759, 2024.
- Y. Song, S. Bhattacharya, and P. Sur. Generalization error of min-norm interpolators in transfer learning. *arXiv preprint arXiv:2406.13944*, 2024.
- G. W. Stewart and J.-g. Sun. Matrix perturbation theory. *Academic Press*, 1990.

- M. Sypetkowski, M. Rezanejad, S. Saberian, O. Kraus, J. Urbanik, J. Taylor, B. Mabey, M. Victors, J. Yosinski, A. R. Sereshkeh, et al. Rxrx1: A dataset for evaluating experimental batch correction methods. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 4285–4294, 2023.
- B. Trabucco, K. Doherty, M. A. Gurinas, and R. Salakhutdinov. Effective data augmentation with diffusion models. In *International Conference on Learning Representations*, 2024.
- B. van Breugel, T. Liu, D. Oglic, and M. van der Schaar. Synthetic data in biomedicine via generative artificial intelligence. *Nature Reviews Bioengineering*, 2(12):991–1004, 2024.
- R. Vershynin. *High-dimensional probability: An introduction with applications in data science*, volume 47. Cambridge university press, 2018.
- Y. Wang, Y. Kordi, S. Mishra, A. Liu, N. A. Smith, D. Khashabi, and H. Hajishirzi. Self-instruct: Aligning language models with self-generated instructions. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 13484–13508, 2023.
- D. Wu and J. Xu. On the optimal weighted ℓ_2 regularization in overparameterized linear regression. In *Advances in Neural Information Processing Systems*, volume 33, pages 10112–10123, 2020.
- S. Wyllie, I. Shumailov, and N. Papernot. Fairness feedback loops: training on synthetic data amplifies bias. In *Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency*, pages 2113–2147, 2024.
- E. Xie, J. Chen, Y. Zhao, J. YU, L. Zhu, Y. Lin, Z. Zhang, M. Li, J. Chen, H. Cai, et al. Sana 1.5: Efficient scaling of training-time and inference-time compute in linear diffusion transformer. In *International Conference on Machine Learning*, 2025.
- F. Yang, H. R. Zhang, S. Wu, C. Re, and W. J. Su. Precise high-dimensional asymptotics for quantifying heterogeneous transfers. *Journal of Machine Learning Research*, 26(113):1–88, 2025.

A Related work

On the theoretical side, we focus on the high-dimensional regime in which both the number of features (i.e., dimension of β) and the number of samples (i.e., dimensions of y_s, y_t) are large and scale proportionally. This setup was considered by a line of research using random matrix theory to characterize test error and various associated phenomena (e.g., benign overfitting [Bartlett et al., 2020] and double descent [Belkin et al., 2019]). More precisely, the test error of ridge(less) regression was studied by Hastie et al. [2022], Wu and Xu [2020], Richards et al. [2021], Cheng and Montanari [2024], the distribution of the ERM solution by Montanari et al. [2019], Chang et al. [2021], Han and Xu [2023], and the impact of spurious correlations by Bombari and Mondelli [2025]. This motivates us to look for practical insights into synthetic data selection by performing a high-dimensional regression analysis. Closer to our work are specific analyses involving more than one distribution, which in our case are the training/test distribution and the synthetic one used for augmentation. More precisely, the test error under distribution shift was analyzed by Patil et al. [2024], Mallinar et al. [2024], but this assumes training on one distribution and testing on the other, as opposed to training on both and testing on one. Training on surrogate data was considered by Ildiz et al. [2025], Kolossov et al. [2024], Jain et al. [2024]: Ildiz et al. [2025] assume that the surrogate data comes from a teacher model and study the phenomenon of weak-to-strong generalization; Kolossov et al. [2024] consider data selection given unlabeled samples plus access to a surrogate model that predicts the labels better than random guessing; Jain et al. [2024] integrate surrogate and real data, but the analysis is limited to isotropic covariance. Most closely related to our theoretical setting is when training occurs on multiple data distributions and testing occurs on a single one of them, which was analyzed both in under-parameterized [Yang et al., 2025] and over-parameterized [Song et al., 2024] regimes. However, Yang et al. [2025], Song et al. [2024] assume that the data distributions have zero mean, which is unrealistic in our context. In fact, centering the data would require access to the mean of the test sample, which is equivalent to having access to its unknown label.

On the practical side, several papers studied how to incorporate synthetic data into training predictors. Besides simply training better generative models, empirical work focused on upgrading the sampling process itself, under the assumption that better conditional generation would lead to more accurate predictors. More precisely, the CLIP model [Radford et al., 2021] underpins many filtering and selection algorithms for generative data. He et al. [2023] propose using CLIP similarity to labels to prune low-quality samples from augmentations. Lin et al. [2023] introduce sampling and filtering strategies based on CLIP similarity to either labels or the mean representation of real data, incorporating diversity via clustering. Almost concurrently, other works argued that synthetic images underperform in scaling laws [Fan et al., 2024] and, if the generative model is pre-trained on external data, simple retrieval baselines can be better [Geng et al., 2024, Burg et al., 2023]. Our work can be interpreted as a more fine-grained investigation of the same problem, characterizing which properties of the generated data improve generalization. At the same time, our results do not preclude that the extra data is real data from another dataset, as tested in Figure 2 in Appendix F. Closer to our solution, Hulkund et al. [2025] explore the problem of data selection given a fixed test set and, taking a purely empirical stance, compare several filtering methods, including an approach inspired by Gadre et al. [2023] that selects clusters of image embeddings. As a heuristic, we find that this works rather well but has shortcomings, as empirically demonstrated in Table 3 in Appendix F.

B Model assumptions

Recall that we work with datasets $X_{(i)}$, for $(i) \in \{t, s\}$, rows of which are independent random vectors with $p \times p$ population covariance matrix $\Sigma_{(i)}$ and mean $\mu_{(i)}$. This can be written as:

$$X_{(i)} = Z^{(i)}(\Sigma_{(i)})^{1/2} + 1_{n_{(i)}}\mu_{(i)}^\top \in \mathbb{R}^{n_{(i)} \times p}, \quad (\text{B.1})$$

where $Z^{(i)} \in \mathbb{R}^{n_{(i)} \times p}$, $\mu_{(i)} \in \mathbb{R}^p$, $1_{n_{(i)}} \in \mathbb{R}^{n_{(i)}}$ is the all-ones vector, and all entries $[Z_{jk}^{(i)}]$ are independent with zero mean and unit variance. By omitting subscripts, we denote by (X, y) the two datasets (X_t, y_t) and (X_s, y_s) stacked, i.e., $X := \begin{bmatrix} X_t \\ X_s \end{bmatrix} \in \mathbb{R}^{n \times p}$, $y := \begin{bmatrix} y_t \\ y_s \end{bmatrix} \in \mathbb{R}^n$.

We make assumptions on the data distribution which are common in related work [Yang et al., 2025, Song et al., 2024]. Let $\tau > 0$ be a small constant. We assume that, for $\psi > 4$, the ψ -th moment of $Z_{jk}^{(i)}$ is upper bounded by $1/\tau$, i.e., $\mathbb{E}[|Z_{jk}^{(i)}|^\psi] \leq \tau^{-1}$, which means that the tails do not decay too

slowly. The eigenvalues of $\Sigma_{(i)}$, denoted as $\lambda_1^{(i)}, \dots, \lambda_p^{(i)}$, are all bounded between τ and τ^{-1} , i.e., $\tau \leq \lambda_p^{(i)} \leq \dots \leq \lambda_2^{(i)} \leq \lambda_1^{(i)} \leq \tau^{-1}$, which means that the covariance matrix is well-conditioned (i.e., the distribution is well-spread). Furthermore, the entries of $\varepsilon_{(i)} \in \mathbb{R}^{n_i}$ have bounded moments up to any order, i.e., for any $k \in \mathbb{N}$, there exists a constant $C_k > 0$ s.t. $\mathbb{E}[|\varepsilon_{(i)j}|^k] \leq C_k$ (noise is not heavy tailed). The sample sizes are comparable with the dimension p , i.e., $\gamma := n/p$, $\gamma_t := n_t/p$, and $\gamma_s := n_s/p$, with $0 \leq \gamma_t \leq 1/\tau$ and $\tau \leq \gamma$, $\gamma_s \leq 1/\tau$. Lastly, let $\|\mu_{(i)}\|_2 = r_{(i)}\sqrt{p}$, where $r_{(i)}$ is a constant, with a constant angle between them $\varphi := |\langle \mu_s, \mu_t \rangle| / (\|\mu_s\|_2 \|\mu_t\|_2)$. This is a technical assumption to simplify the proof notation. If φ is allowed to depend on n, p , all results (and corresponding proofs) still hold verbatim, as long as either $\varphi < 1 - \delta$ for some constant $\delta > 0$ or $\varphi = 1$.

C Under-parametrized regime addendum

C.1 Training only on synthetic data

We adjust the assumption at the beginning of Section 3. Namely, we assume that $\gamma_t = 0$, $1 + \tau \leq \gamma_s = \gamma \leq 1/\tau$, which means that we are training on data from a single distribution that is different from the one we are testing on.

Proposition C.1. *In the setting described above, it holds that, with high probability,*

$$\lim_{n \rightarrow \infty} \left| R_X(\hat{\beta}; \beta) - \frac{\sigma^2}{n} \cdot \frac{\gamma}{\gamma - 1} \cdot \left[\text{Tr}[\Sigma_t \Sigma_s^{-1}] + \|\Sigma_s^{-1/2} \mu_t\|_2^2 - \left(\frac{\mu_t^\top \Sigma_s^{-1} \mu_s}{\|\Sigma_s^{-1/2} \mu_s\|_2} \right)^2 \right] \right| = 0. \quad (\text{C.1})$$

This result (proved in Appendix E.3) extends the zero-centered expression by Hastie et al. [2022]. We observe consistency if we disregard means ($\mu_s = \mu_t = 0$) and covariance shift ($\Sigma_t \Sigma_s^{-1} = I_p$). Proposition C.1 also extends the zero-centered anisotropic setting of Yang et al. [2025] to the case without samples from the training distribution, and consistency follows after setting $\mu_s = \mu_t = 0$. The effect of the mean shift is captured by $\|\Sigma_s^{-1/2} \mu_t\|_2^2 - (\mu_t^\top \Sigma_s^{-1} \mu_s / \|\Sigma_s^{-1/2} \mu_s\|_2)^2$: what matters is (i) the cosine similarity between $\Sigma_s^{-1/2} \mu_s$ and $\Sigma_s^{-1/2} \mu_t$, and (ii) the alignment of the principal directions of Σ_s with μ_t . In other words, the excess risk decreases as (i) the mean of synthetic training data aligns with the mean of test data in the directions of the synthetic covariance matrix, and (ii) the principal directions of the synthetic covariance matrix align with the test mean.

C.2 Trace normalization

Note that, increasing the scale of Σ_s also reduces the risk proxy $\mathcal{R}_u(\Sigma_t^{-1/2} \Sigma_s \Sigma_t^{-1/2})$, which we formulate in the following proposition. See Appendix E.5 for the proof and Figure 1c for an illustration.

Proposition C.2. *For any $M \in \mathbb{R}_{>0}^{p \times p}$ and any constant $\eta > 1$, it holds that $\mathcal{R}_u(\eta M) \leq \mathcal{R}_u(M)$.*

Recalling $M = \Sigma_t^{-1/2} \Sigma_s \Sigma_t^{-1/2}$, this suggests that greater diversity in synthetic data is advantageous. However, as Theorem 3.1 relies on bounds on the spectra of Σ_t, Σ_s (see Section 2), η must be of constant order, i.e., it cannot grow with n and p (otherwise, the error between $R_X(\hat{\beta}; \beta)$ and $\mathcal{R}_u(\eta M)$ may not vanish as in (3.1)). This motivates the trace normalization ($\text{Tr}[M] = p$) in Theorem 3.2. While other normalizations exist (e.g., on the determinant in [Yang et al., 2025]), they overly restrict the search space and make interpretation for synthetic data selection less clear.

C.3 Numerical simulation

We illustrate the theoretical insights from Section 3 via synthetic numerical experiments. Figure 1a highlights the independence of the test error on the mean shift, as suggested by Theorem 3.1. Namely, it shows that the excess risk remains unchanged upon varying the cosine similarity between the means. Moreover, the conclusion around Theorem 3.2 is corroborated by Figure 1b, illustrating that the risk decreases as the covariance matrices Σ_t and Σ_s become more similar.

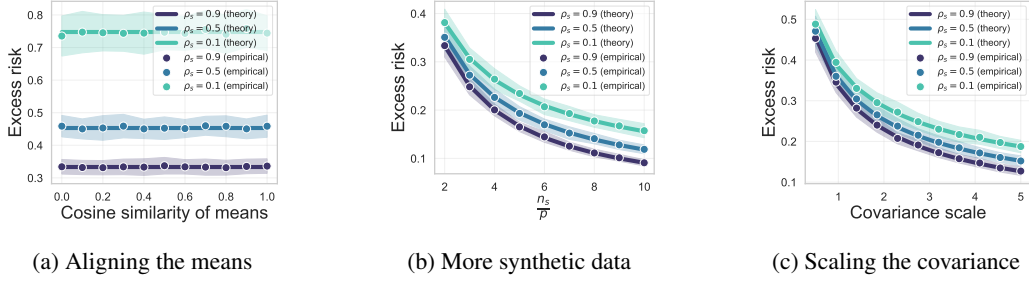


Figure 1: Excess risk using training data from $\mathcal{N}(\mu_t, \Sigma_t)$ and synthetic data from $\mathcal{N}(\mu_s, \Sigma_s)$, where Σ_t, Σ_s are Kac–Murdock–Szegő matrices (Toeplitz matrices with geometrically decaying entries) with parameters ρ_t, ρ_s , scaled so that $\text{Tr}[M] = p$. We pick $\|\mu_t\|_2 = \|\mu_s\|_2 = 2\sqrt{p}$, $\rho_t = 0.9$, $p = 600$, $n_t = 1200$, $n_s = 1200$, unless varying the parameters in the plot. Each value is computed from 100 i.i.d. trials, the error band is at 1 standard deviation, and theoretical predictions are continuous lines. Different curves correspond to different values of ρ_s . (a) Changing the cosine similarity of the mean does not impact the risk (here, Σ_s is scaled by $\eta := \rho_s$). (b) Larger ρ_s gives lower risk since Σ_s is closer to Σ_t . (c) Scaling Σ_s reduces the risk.

D Over-parameterized regime

As opposed to Section 3, let us assume that $\tau \leq \gamma, \gamma_s, \gamma_t \leq 1/(1 + \tau)$, so that $n < p$ and we are in the over-parameterized regime. We sample β from a sphere of constant radius, independently from $X, \varepsilon_t, \varepsilon_s$. We also assume that Σ_s and Σ_t are simultaneously diagonalizable. This assumption is of technical nature and common in related work [Song et al., 2024, Mallinar et al., 2024, Ildiz et al., 2025]. Writing out this condition, we have the SVDs $\Sigma_s = U\Lambda^s U^\top, \Sigma_t = U\Lambda^t U^\top$. Let us denote by $\lambda_i^s := \Lambda_{i,i}^s, \lambda_i^t := \Lambda_{i,i}^t$ and introduce the spectral probability distributions used in our claims:

$$\hat{H}_p(\lambda^s, \lambda^t) := \frac{1}{p} \sum_{i=1}^p \mathbf{1}_{\{(\lambda^s, \lambda^t) = (\lambda_i^s, \lambda_i^t)\}}, \quad \hat{G}_p(\lambda^s, \lambda^t) := \sum_{i=1}^p \langle \beta, u_i \rangle^2 \mathbf{1}_{\{(\lambda^s, \lambda^t) = (\lambda_i^s, \lambda_i^t)\}}. \quad (\text{D.1})$$

This section follows the same blueprint as Section 3 for the under-parameterized regime. Namely, Theorem D.1 gives a *deterministic equivalent* of the excess risk using training and synthetic data and, in doing so, it extends results by Song et al. [2024] to non-zero centered data. The deterministic equivalent depends only on regression coefficients β and covariances Σ_t, Σ_s , and it does not depend on means μ_t, μ_s . Then, Theorem D.2 finds Σ_s that minimizes the limit risk from Theorem D.1 when $\Sigma_t = I_p$, thus showing the optimality of *covariance matching* ($\Sigma_s \propto \Sigma_t$) with isotropic training data. The proofs of these results follow a similar argument chain as in Section 3, although they tend to be more technically involved. We briefly discuss differences, deferring the full arguments of Theorems D.1 and D.2 to Appendices E.7 and E.8, respectively.

Theorem D.1. *Under the assumptions from Section 2 and the start of this section, it holds that, with high probability,*

$$\lim_{n \rightarrow \infty} \left| R_X(\hat{\beta}; \beta) - \mathcal{V}(\Sigma_s, \Sigma_t) - \mathcal{B}(\Sigma_s, \Sigma_t, \beta) \right| = 0, \quad (\text{D.2})$$

where

$$\mathcal{V}(\Sigma_s, \Sigma_t) := \frac{\sigma^2}{\gamma} \int \frac{-\lambda^t (a_3 \lambda^s + a_4 \lambda^t)}{(a_1 \lambda^s + a_2 \lambda^t + 1)^2} d\hat{H}_p(\lambda^s, \lambda^t), \quad \mathcal{B}(\Sigma_s, \Sigma_t, \beta) := \int \frac{b_3 \lambda^s + (b_4 + 1) \lambda^t}{(b_1 \lambda^s + b_2 \lambda^t + 1)^2} d\hat{G}_p(\lambda^s, \lambda^t),$$

and a_i, b_i ($i \in \{1, 2, 3, 4\}$) are the unique solutions to the equations reported in Appendix E.6.

We highlight two additional difficulties in the proof of Theorem D.1 arising from the over-parameterized regime: (1) the inverse does not replace the pseudo-inverse in (E.2), and (2) the bias term does not vanish. We address the former by introducing the λ -regularized ridge estimator $\hat{\beta}_\lambda$, which approximates $\hat{\beta}$ for small λ and admits inverse-based formulas similar to (E.2). Addressing the latter requires a delicate control of the inverse, obtained via Woodbury formula.

Theorem D.2. *Let $\mathcal{S} := \{\Sigma \in \mathbb{R}_{>0}^{p \times p} : \text{Tr}(\Sigma) = p\}$, where $\mathbb{R}_{>0}^{p \times p}$ denotes the set of $p \times p$ positive definite matrices. Recall the definitions of $\mathcal{V}(\Sigma_s, \Sigma_t)$, $\mathcal{B}(\Sigma_s, \Sigma_t, \beta)$ from Theorem D.1, and define $\mathcal{R}_o(\Sigma_s, \Sigma_t, \beta) := \mathcal{V}(\Sigma_s, \Sigma_t) + \mathcal{B}(\Sigma_s, \Sigma_t, \beta)$. Then, for any $\Sigma_s \in \mathcal{S}$, with high probability over the sampling of β over a sphere of constant radius, it holds that*

$$\mathcal{R}_o(I_p, I_p, \beta) \leq \mathcal{R}_o(\Sigma_s, I_p, \beta) + o(1),$$

where $o(1)$ denotes a quantity that vanishes as $n, p \rightarrow \infty$.

Due to the complexity of the expressions for $\mathcal{V}(\Sigma_s, \Sigma_t)$ and $\mathcal{B}(\Sigma_s, \Sigma_t, \beta)$, the optimality of covariance matching ($\Sigma_s \propto \Sigma_t$) in the over-parameterized regime is shown for isotropic training data ($\Sigma_t = I_p$). At the technical level, we note that the bias generally depends on the eigenspace decomposition of the covariance matrices via $\hat{G}_p$, as defined in (D.1). However, when $\Sigma_t = I_p$, cancellations in the equations for b_i ($i \in \{1, 2, 3, 4\}$) give that the bias $\mathcal{B}(\Sigma_s, I_p, \beta)$ is close to $\frac{p-n}{p} \|\beta\|_2$ for any Σ_s . Having obtained that, the variance is then optimized following the approach of Theorem 3.2.

E Proofs of the theoretical results

Additional notation. We use the shorthand $[n] := \{1, \dots, n\}$ for an integer n . Given a matrix M , its operator norm is denoted by $\|M\|_2$, its i -th largest singular value by $\sigma_i(M)$ and the corresponding i -th left-singular (resp. right-singular) vector of unit norm by $u_i(M)$ (resp. $v_i(M)$). Additionally, when applicable, we denote the i -th largest eigenvalue of M by $\lambda_i(M)$. We use $\mathbb{R}_{>0}^{p \times p}$ to denote the set of all $p \times p$ positive definite matrices, and S^{p-1} to denote a $(p-1)$ -dimensional unit sphere. We denote by e_i the i -th element of the canonical basis of $\mathbb{R}^l$, where the exact exponent l is assumed from context. We will say that an event $\mathcal{E}$ happens *with high probability* (w.h.p.) if and only if $\mathbb{P}(\mathcal{E}) \rightarrow 1$ as $p, n \rightarrow \infty$. Moreover, we will say that an event Ξ happens *with overwhelming probability* if and only if, for any large constant $D > 0$, $\mathbb{P}(\Xi) \geq 1 - p^{-D}$ for large enough p . Lastly, throughout this appendix, we use c to denote a constant (independent of n, p) whose value may change from line to line.

For convenience, we recall some notation and definitions from Section 2. Namely, we denote by $Z \in \mathbb{R}^{n \times p}$ a random matrix with i.i.d. entries having zero mean, unit variance and bounded ψ -th moment (for some $\psi > 4$). Recall $\mu_{(i)} \in \mathbb{R}^p$, for $(i) \in \{s, t\}$, such that $\|\mu_{(i)}\|_2 = r_{(i)}\sqrt{p}$, where $r_{(i)}$ is a constant, with a constant angle between them $\varphi := |\langle \mu_s, \mu_t \rangle| / (\|\mu_s\|_2 \|\mu_t\|_2)$. Also, let $\Sigma_s, \Sigma_t \in \mathbb{R}^{p \times p}$ be covariance matrices with bounds on their spectrum as in Section 2. Then, we consider

a data distribution $X = \begin{bmatrix} Z_t \Sigma_t^{1/2} + 1_{n_t} \mu_t^\top \\ Z_s \Sigma_s^{1/2} + 1_{n_s} \mu_s^\top \end{bmatrix} \in \mathbb{R}^{n \times p}$ and introduce its zero mean counterpart

$X^0 := \begin{bmatrix} Z_t \Sigma_t^{1/2} \\ Z_s \Sigma_s^{1/2} \end{bmatrix}$. The corresponding sample covariance matrices are defined as $\hat{\Sigma} = \frac{X^\top X}{n}$ and $\hat{\Sigma}_0 = \frac{X^{0\top} X^0}{n}$. Lastly, unless stated otherwise, we work in the regime $n/p = \gamma$, where $\gamma \neq 1$ is a fixed constant independent of n and p .

E.1 Bias and variance decomposition

The excess risk as defined in (2.2) can be further decomposed into bias and variance as

$$R_X(\hat{\beta}; \beta) = \|\mathbb{E}[\hat{\beta} | X] - \beta\|_{\Sigma_t + \mu_t \mu_t^\top}^2 + \text{Tr}[\text{Cov}(\hat{\beta} | X)(\Sigma_t + \mu_t \mu_t^\top)] = B_X(\hat{\beta}; \beta) + V_X(\hat{\beta}; \beta). \quad (\text{E.1})$$

Then, plugging the known expression for the min norm interpolator $\hat{\beta} = (X^\top X)^+ X^\top y$ (see e.g. [Yang et al., 2025, Equation (2.7)]) into (E.1) yields closed-form expression for bias and variance, i.e.,

$$B_X(\hat{\beta}; \beta) = \beta^\top \Pi (\Sigma_t + \mu_t \mu_t^\top) \Pi \beta \quad \text{and} \quad V_X(\hat{\beta}; \beta) = \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}^+ (\Sigma_t + \mu_t \mu_t^\top)], \quad (\text{E.2})$$

where $\hat{\Sigma} = X^\top X/n$, and $\Pi = I - \hat{\Sigma}^+ \hat{\Sigma}$ (projection on the null space of X).

E.2 Proof of Theorem 3.1

We first state and prove useful results, in which we analyze the behavior of singular values of a low-rank perturbation of matrices.

Proposition E.1. Let $\sigma_1 \geq \dots \geq \sigma_{\min(n,p)}$ be the singular values of $\tilde{Z} = \frac{Z + 1_n \mu_s^\top}{\sqrt{n}}$. Then, there exists a constant $c(\gamma) > 0$ independent of n , such that, almost surely,

$$\liminf_{n \rightarrow \infty} \sigma_{\min(n,p)} \geq c(\gamma).$$

Proof. To simplify notation we will refer to $\sigma_{\min}$ as the smallest singular value of a matrix. Let us choose an orthogonal matrix $Q \in \mathbb{R}^{n \times n}$ such that $Q1_n = \sqrt{n} e_1$. Since singular values are left orthogonally invariant, we may replace $\tilde{Z}$ by

$$\tilde{Z}' = \frac{QZ}{\sqrt{n}} + e_1 \mu_s^\top.$$

Writing the rows of QZ as

$$QZ = \begin{bmatrix} z_1^\top \\ Z_2 \end{bmatrix}, \quad z_1 \in \mathbb{R}^p, \quad Z_2 \in \mathbb{R}^{(n-1) \times p},$$

we have

$$\tilde{Z}' = \begin{bmatrix} \frac{z_1^\top}{\sqrt{n}} + \mu_s^\top \\ \frac{Z_2}{\sqrt{n}} \end{bmatrix}.$$

For any unit vector $x \in \mathbb{R}^p$,

$$\|\tilde{Z}'x\|_2 = \sqrt{\left(\frac{z_1^\top x}{\sqrt{n}} + \mu_s^\top x\right)^2 + \left\|\frac{Z_2 x}{\sqrt{n}}\right\|_2^2} \geq \left\|\frac{Z_2 x}{\sqrt{n}}\right\|_2.$$

Hence, by the variational definition of singular values, we have

$$\sigma_{\min}(\tilde{Z}) \geq \sigma_{\min}\left(\frac{Z_2}{\sqrt{n}}\right) = \sqrt{\frac{n-1}{n}} \sigma_{\min}\left(\frac{Z_2}{\sqrt{n-1}}\right). \quad (\text{E.3})$$

By the Bai–Yin theorem [Bai and Silverstein, 2010, Theorem 5.11], for an $(n-1) \times p$ random matrix Z_2 with i.i.d entries with mean zero, unit variance and bounded fourth moments it holds

$$\sigma_{\min}\left(\frac{Z_2}{\sqrt{n-1}}\right) \xrightarrow[n \rightarrow \infty]{a.s.} |1 - \sqrt{p/(n-1)}|.$$

Therefore, applying $\liminf_{n \rightarrow \infty}$ to (E.3), we have

$$\begin{aligned} \liminf_{n \rightarrow \infty} \sigma_{\min}(\tilde{Z}) &\geq \liminf_{n \rightarrow \infty} \sqrt{\frac{n-1}{n}} \sigma_{\min}\left(\frac{Z_2}{\sqrt{n}}\right) \\ &= \lim_{n \rightarrow \infty} \sqrt{\frac{n-1}{n}} \sigma_{\min}\left(\frac{Z_2}{\sqrt{n}}\right) \\ &= |1 - \gamma^{-1/2}| > 0, \end{aligned}$$

which gives the desired result as $\gamma \neq 1$. □

Proposition E.2. Let $\tilde{X}_n = X/\sqrt{n} = \frac{1}{\sqrt{n}} \begin{bmatrix} Z_t \Sigma_t^{1/2} + 1_{n_t} \mu_t^\top \\ Z_s \Sigma_s^{1/2} + 1_{n_s} \mu_s^\top \end{bmatrix} \in \mathbb{R}^{n \times p}$. Let $\sigma_1 \geq \dots \geq \sigma_n$ be the singular values of $\tilde{X}_n$ and $v_1, \dots, v_n$ be the corresponding right singular vectors. Then, as $n \rightarrow \infty$, the following results hold:

1. For $\varphi < 1$, we have

1a. $\sigma_1 = \Theta(\sqrt{p})$, $\sigma_2 = \Theta(\sqrt{p})$, and $\sigma_3 = O(1)$;

$$\begin{aligned} 1b. \quad &\left| \frac{\langle v_1, \mu_s \rangle}{\|v_1\|_2 \|\mu_s\|_2} \right|^2 + \left| \frac{\langle v_2, \mu_s \rangle}{\|v_2\|_2 \|\mu_s\|_2} \right|^2 = 1 - O\left(\frac{1}{p}\right), \\ &\left| \frac{\langle v_1, \mu_t \rangle}{\|v_1\|_2 \|\mu_t\|_2} \right|^2 + \left| \frac{\langle v_2, \mu_t \rangle}{\|v_2\|_2 \|\mu_t\|_2} \right|^2 = 1 - O\left(\frac{1}{p}\right). \end{aligned}$$

2. For $\varphi = 1$, we have

2a. $\sigma_1 = \Theta(\sqrt{p})$, $\sigma_2 = O(1)$;

$$2b. \quad \left| \frac{\langle v_1, \mu_s \rangle}{\|v_1\|_2 \|\mu_s\|_2} \right|^2 = 1 - O\left(\frac{1}{p}\right).$$

Proof. Let us first abuse notation and write $1_{n_s} = [1, \dots, 1, 0, \dots, 0]^\top \in \mathbb{R}^{n \times 1}$ (n_s ones followed by n_t zeros) and $1_{n_t} = [0, \dots, 0, 1, \dots, 1]^\top \in \mathbb{R}^{n \times 1}$ (n_s zeros followed by n_t ones). Then if we write

$$X_n := \frac{1}{\sqrt{n}} \begin{bmatrix} Z_t \Sigma_t^{1/2} \\ Z_s \Sigma_s^{1/2} \end{bmatrix},$$

it holds

$$\tilde{X}_n = X_n + P_n, \quad \text{where} \quad P_n := \frac{1_{n_s} \mu_s^\top + 1_{n_t} \mu_t^\top}{\sqrt{n}}. \quad (\text{E.4})$$

To obtain the wanted result, we will need to express the non-zero singular values and the corresponding right singular vectors of the rank-2 perturbation P_n , that is $\sigma_i(P_n)$ and $v_i(P_n)$, $i \in [2]$. Notice that

$$P_n^\top P_n = \alpha_s^2 \mu_s \mu_s^\top + \alpha_t^2 \mu_t \mu_t^\top,$$

where $\alpha_s := \sqrt{\frac{n_s}{n}}$ and $\alpha_t := \sqrt{\frac{n_t}{n}}$. Moreover, it holds

$$P_n^\top P_n = Q_p^\top Q_p, \quad (\text{E.5})$$

where $Q_p = \begin{bmatrix} \alpha_s \mu_s \\ \alpha_t \mu_t \end{bmatrix} \in \mathbb{R}^{2 \times p}$. Note that

$$Q_p Q_p^\top = \begin{bmatrix} \alpha_s^2 \|\mu_s\|_2^2 & \alpha_s \alpha_t \langle \mu_s, \mu_t \rangle \\ \alpha_s \alpha_t \langle \mu_s, \mu_t \rangle & \alpha_t^2 \|\mu_t\|_2^2 \end{bmatrix} =: \begin{bmatrix} a & b \\ b & d \end{bmatrix},$$

and it is enough to analyze its SVD, since

$$\sigma_i(P_n) = \sqrt{\sigma_i(Q_p Q_p^\top)}, \quad \text{and} \quad v_i(P_n) = \frac{1}{\sigma_i(P_n)} v_i(Q_p Q_p^\top)^\top Q_p.$$

The previous equations hold due to (E.5), since $\sigma_i(Q_p) = \sigma_i(P_n)$, and

$$\frac{1}{\sigma_i(P_n)} v_i(Q_p Q_p^\top)^\top Q_p = \frac{1}{\sigma_i(Q_p)} u_i(Q_p)^\top [u_1(Q_p) \quad u_2(Q_p)] \begin{bmatrix} \sigma_1(Q_p) & 0 \\ 0 & \sigma_2(Q_p) \end{bmatrix} \begin{bmatrix} v_1(Q_p) \\ v_2(Q_p) \end{bmatrix}.$$

This implies that, for $i \in [2]$, the singular vectors $v_i(P_n)$ are in the span $\{\mu_s, \mu_t\}$. Recall that the angle between μ_s and μ_t is fixed to $\varphi := \frac{|\langle \mu_s, \mu_t \rangle|}{\|\mu_s\|_2 \|\mu_t\|_2}$.

We first consider the case when $\varphi < 1$. It holds that the eigenvalues of $Q_p Q_p^\top$ are

$$\begin{aligned} \sigma_{1,2}(Q_p Q_p^\top) &= \frac{a + d \pm \sqrt{(a - d)^2 + 4b^2}}{2} \\ &= \frac{(r_s^2 \alpha_s^2 + r_t^2 \alpha_t^2) p \pm \sqrt{(r_s^2 \alpha_s^2 - r_t^2 \alpha_t^2)^2 p^2 + 4 \alpha_s^2 r_s^2 \alpha_t^2 r_t^2 \varphi^2 p^2}}{2} \\ &\geq p \cdot \frac{(r_s^2 \alpha_s^2 + r_t^2 \alpha_t^2) - \sqrt{(r_s^2 \alpha_s^2 - r_t^2 \alpha_t^2)^2 + 4 \alpha_s^2 r_s^2 \alpha_t^2 r_t^2 \varphi^2}}{2} \\ &= p \cdot c_1, \end{aligned} \quad (\text{E.6})$$

with $c_1 = \frac{(r_s^2 \alpha_s^2 + r_t^2 \alpha_t^2) - \sqrt{(r_s^2 \alpha_s^2 - r_t^2 \alpha_t^2)^2 + 4 \alpha_s^2 r_s^2 \alpha_t^2 r_t^2 \varphi^2}}{2} > 0$, since $\varphi < 1$. This implies that

$$\sigma_i(P_n) \geq c \cdot \sqrt{p},$$

for some constant c .

Furthermore, it almost surely holds that

$$\begin{aligned} \sigma_1(X_n) &= \sqrt{\sigma_1(X_n^\top X_n)} \\ &= \sqrt{\sigma_1(\Sigma_s^{1/2} Z_s^\top Z_s \Sigma_s^{1/2} + \Sigma_t^{1/2} Z_t^\top Z_t \Sigma_t^{1/2})} \\ &\leq \sqrt{\sigma_1(\Sigma_s^{1/2} Z_s^\top Z_s \Sigma_s^{1/2})} + \sigma_1(\Sigma_t^{1/2} Z_t^\top Z_t \Sigma_t^{1/2}) \\ &\leq \sqrt{2(1 + \sqrt{\gamma})^2 \cdot \tau^{-1}} = O(1), \end{aligned}$$

due to the convergences of the largest eigenvalue of the sample covariance matrices $Z_s^\top Z_s$ and $Z_t^\top Z_t$ by Bai–Yin theorem [Bai and Silverstein, 2010, Theorem 5.11] and the boundedness of the spectrum

of Σ_s and Σ_t . Then, from Weyl's inequality for singular values (see e.g. [Horn and Johnson, 2012, Chapter 7]), we have that

$$\begin{aligned}\sigma_i(X_n + P_n) &\geq \sigma_i(P_n) - \sigma_1(X_n), \quad \text{for } i = 1, 2, \\ \sigma_3(X_n + P_n) &\leq \sigma_3(P_n) + \sigma_1(X_n) = \sigma_1(X_n),\end{aligned}$$

which implies that $\sigma_{1,2}(X_n + P_n) \geq c \cdot \sqrt{p}$, whereas $\sigma_i(X_n + P_n) = O(1)$, for $i \geq 3$. For the upper bound, note that from (E.6) it holds

$$\sigma_{1,2}(QQ^\top) \leq \frac{(r_s^2 \alpha_s^2 + r_t^2 \alpha_t^2)p + (r_s^2 \alpha_s^2 + r_t^2 \alpha_t^2)p}{2} = p \cdot c_2,$$

implying $\sigma_i(P_n) \leq c \cdot \sqrt{p}$. Applying Weyl's inequality for singular values once more, we get

$$\sigma_i(X_n + P_n) \leq \sigma_1(X_n) + \sigma_i(P_n) = O(\sqrt{p}),$$

concluding the proof of **1a**.

Moving onto singular vectors, let us recall the definition of spectral distance between two k -dimensional subspaces $\mathcal{W} \leq \mathbb{R}^p$ and $\tilde{\mathcal{W}} \leq \mathbb{R}^p$, as it will be used to conclude the proof. Towards this end, we first introduce principal angles $\theta_1 \dots \theta_k \in [0, \pi/2]$ between $\mathcal{W}$ and $\tilde{\mathcal{W}}$, which are defined recursively from $i = 1$ as

$$\cos(\theta_i) = \max_{w_i \in \mathcal{W}, \tilde{w}_i \in \tilde{\mathcal{W}}} \frac{\langle w_i, \tilde{w}_i \rangle}{\|w_i\|_2 \|\tilde{w}_i\|_2},$$

subject to $w_i, \tilde{w}_i$ being orthogonal to the previous maximizers. Then, the spectral distance between $\mathcal{W}$ and $\tilde{\mathcal{W}}$ is defined as

$$d(\mathcal{W}, \tilde{\mathcal{W}}) := \max_{i \in [k]} \sin \theta_i.$$

There is an alternative way to express this spectral distance between subspaces, using their orthonormal basis. Namely, let $W \in \mathbb{R}^{p \times k}$ and $\tilde{W} \in \mathbb{R}^{p \times k}$ be such that their columns form an orthonormal basis of $\mathcal{W}$ and $\tilde{\mathcal{W}}$, respectively. Then by [Stewart and Sun, 1990, Chapter II, Corollary 5.4] it holds

$$d(\mathcal{W}, \tilde{\mathcal{W}}) := \left\| (I - WW^\top) \tilde{W} \right\|_2. \quad (\text{E.7})$$

Let us denote by $\tilde{V} := \begin{bmatrix} v_1(P_n) \\ v_2(P_n) \end{bmatrix}^\top$, $V := \begin{bmatrix} v_1(\tilde{X}_n) \\ v_2(\tilde{X}_n) \end{bmatrix}^\top$ and by $\mathcal{V}$, $\tilde{\mathcal{V}}$ the subspaces spanned by their columns. Then, by Wedin's $\sin \Theta$ theorem, [Stewart and Sun, 1990, Chapter V, Theorem 4.4.] it holds that

$$d(\mathcal{V}, \tilde{\mathcal{V}}) \leq \frac{\sigma_1(X_n)}{\sigma_2(X_n + P_n) - \sigma_3(X_n + P_n)} = \frac{1}{c \cdot \sqrt{p} + O(1)} = O\left(\frac{1}{\sqrt{p}}\right).$$

As $v_1(P_n), v_2(P_n) \in \text{span}\{\mu_s, \mu_t\}$ and they are linearly independent, this implies that $\mathcal{V} = \text{span}\{\mu_s, \mu_t\}$. Choosing matrices $\tilde{V}_s \in \mathbb{R}^{p \times 2}$ and $\tilde{V}_t \in \mathbb{R}^{p \times 2}$ such that their columns are orthonormal bases of $\tilde{\mathcal{V}}$ and their first column is $\frac{\mu_s}{\|\mu_s\|_2}$ and $\frac{\mu_t}{\|\mu_t\|_2}$ respectively, one gets that

$$\begin{aligned}\left\| (I - VV^\top) \frac{\mu_s}{\|\mu_s\|_2} \right\|_2 &= \left\| (I - VV^\top) \tilde{V}_s e_1 \right\|_2 \leq \left\| (I - VV^\top) \tilde{V}_s \right\|_2 = d(\mathcal{V}, \tilde{\mathcal{V}}) \leq O\left(\frac{1}{\sqrt{p}}\right), \\ \left\| (I - VV^\top) \frac{\mu_t}{\|\mu_t\|_2} \right\|_2 &= \left\| (I - VV^\top) \tilde{V}_t e_1 \right\|_2 \leq \left\| (I - VV^\top) \tilde{V}_t \right\|_2 = d(\mathcal{V}, \tilde{\mathcal{V}}) \leq O\left(\frac{1}{\sqrt{p}}\right).\end{aligned}$$

From this, **1b** directly follows. The case $\varphi = 1$ is handled analogously. $\square$

Proposition E.3. *In the under-parameterized regime, i.e., when $p < n$, it holds that*

$$\frac{1}{n} \text{Tr}[\hat{\Sigma}^+(\Sigma_t + \mu_t \mu_t^\top)] = \frac{1}{n} \text{Tr}[\hat{\Sigma}_0^+ \Sigma_t] + O\left(\frac{1}{p}\right). \quad (\text{E.8})$$

Proof. We break down the LHS of (E.8) into two terms

$$\frac{1}{n} \text{Tr}[\hat{\Sigma}^+(\Sigma_t + \mu_t \mu_t^\top)] = T_1 + T_2,$$

where

$$T_1 = \frac{1}{n} \text{Tr}[\hat{\Sigma}^+ \Sigma_t], \quad \text{and} \quad T_2 = \frac{1}{n} \text{Tr}[\hat{\Sigma}^+ \mu_t \mu_t^\top].$$

We will deal with each of the terms individually.

Bounding the term T_1 . It holds that

$$\begin{aligned}
T_1 &= \frac{1}{n} \text{Tr}(\hat{\Sigma}^+ \Sigma_t) \\
&= \frac{1}{n} \text{Tr} \left(\left(\left(\frac{X \Sigma_t^{-1/2}}{\sqrt{n}} \right)^\top \frac{X \Sigma_t^{-1/2}}{\sqrt{n}} \right)^{-1} \right) \\
&= \frac{1}{n} \text{Tr} \left((\bar{X}^\top \bar{X})^{-1} \right) \\
&= \frac{1}{n} \sum_{i=1}^k \frac{1}{\sigma_i^2(\bar{X})},
\end{aligned} \tag{E.9}$$

where $\bar{X} := \frac{X \Sigma_t^{-1/2}}{\sqrt{n}} \in \mathbb{R}^{n \times p}$, and $k \leq p$ is the number of non-zero singular values of $\bar{X}$. Let us prove that $\sigma_p(\bar{X}) > c$ for some constant c , implying that $k = p$. Towards this end, we write out

$$\bar{X}^\top \bar{X} = \bar{X}_s^\top \bar{X}_s + \bar{X}_t^\top \bar{X}_t,$$

where $\bar{X}_s := \frac{X_s \Sigma_t^{-1/2}}{\sqrt{n}} = \frac{(Z_s \Sigma_s^{1/2} + 1_{n_s} \mu_s^\top) \Sigma_t^{-1/2}}{\sqrt{n}}$ and $\bar{X}_t := \frac{X_t \Sigma_t^{-1/2}}{\sqrt{n}} = \frac{(Z_t \Sigma_t^{1/2} + 1_{n_t} \mu_t^\top) \Sigma_t^{-1/2}}{\sqrt{n}}$. From Proposition E.1, it follows that for large enough n , almost surely

$$\sigma_p(\bar{X}_s) \geq c, \quad \sigma_p(\bar{X}_t) \geq c,$$

for some constant c , which is just $c(\gamma)$ from the proposition adjusted by the bound on the eigenvalues of $\Sigma_t^{-1/2}$ and $\Sigma_s^{-1/2}$ (recall that the smallest eigenvalue of Σ_s, Σ_t is lower bounded by τ). Plugging this in gives

$$\sigma_p(\bar{X})^2 \geq \sigma_p(\bar{X}_s)^2 + \sigma_p(\bar{X}_t)^2 \geq 2c^2. \tag{E.10}$$

Let $\bar{X}^0 := \frac{X^0 \Sigma_t^{-1/2}}{\sqrt{n}}$ and note that $\bar{X}$ is a rank-2 perturbation of $\bar{X}^0$ (see (E.4)). Then, due to Weyl's inequality for singular values, it holds that, for $i \in \{3, \dots, p-2\}$,

$$\sigma_{i+2}(\bar{X}^0) \leq \sigma_i(\bar{X}) \leq \sigma_{i-2}(\bar{X}^0).$$

Therefore, we have

$$\frac{1}{n} \sum_{i=1}^{p-4} \frac{1}{\sigma_i(\bar{X}^0)^2} \leq \frac{1}{n} \sum_{i=3}^{p-2} \frac{1}{\sigma_i(\bar{X})^2} \leq \frac{1}{n} \sum_{i=3}^p \frac{1}{\sigma_i(\bar{X}^0)^2}.$$

An application of the Bai–Yin theorem [Bai and Silverstein, 2010, Theorem 5.11] gives that there exist constants a and b such that

$$0 < a < \sigma_p(\bar{X}^0) \leq \sigma_1(\bar{X}^0) < b < +\infty,$$

for large enough n . Therefore, it holds

$$\frac{1}{n} \sum_{i=1}^p \frac{1}{\sigma_i(\bar{X}^0)^2} - O\left(\frac{1}{n}\right) \leq \frac{1}{n} \sum_{i=3}^{p-2} \frac{1}{\sigma_i(\bar{X})^2} \leq \frac{1}{n} \sum_{i=1}^p \frac{1}{\sigma_i(\bar{X}^0)^2},$$

which implies that

$$\frac{1}{n} \sum_{i=3}^{p-2} \frac{1}{\sigma_i(\bar{X})^2} = \frac{1}{n} \sum_{i=1}^p \frac{1}{\sigma_i(\bar{X}^0)^2} + \Theta\left(\frac{1}{n}\right).$$

Using the proved fact that $\sigma_i(\bar{X}) > c$ we have

$$\frac{1}{n} \sum_{i=1}^p \frac{1}{\sigma_i(\bar{X})^2} = \frac{1}{n} \sum_{i=3}^{p-2} \frac{1}{\sigma_i(\bar{X})^2} + O\left(\frac{1}{n}\right).$$

Combining all the pieces, it holds that

$$\begin{aligned}
T_1 &= \frac{1}{n} \sum_{i=1}^p \frac{1}{\sigma_i^2(\bar{X})} \\
&= \frac{1}{n} \sum_{i=3}^{p-2} \frac{1}{\sigma_i(\bar{X})^2} + O\left(\frac{1}{n}\right) \\
&= \frac{1}{n} \sum_{i=1}^p \frac{1}{\sigma_i(\bar{X}^0)^2} + O\left(\frac{1}{n}\right) \\
&= \frac{1}{n} \text{Tr}\left((\bar{X}^{0\top} \bar{X}^0)^{-1} \Sigma_t\right) + O\left(\frac{1}{n}\right) \\
&= \frac{1}{n} \text{Tr}[\hat{\Sigma}_0^+ \Sigma_t] + O\left(\frac{1}{p}\right).
\end{aligned}$$

Bounding the term T_2 . First, recall the shorthand $\tilde{X}_n = X/\sqrt{n}$ and note that

$$\sigma_p(\tilde{X}_n) = \sigma_p(\bar{X} \Sigma_t^{1/2}) \geq \sigma_p(\bar{X}) \cdot \sigma_p(\Sigma_t^{1/2}) \geq c \cdot \tau, \quad (\text{E.11})$$

where the last inequality follows from (E.10) and the bounds on the spectrum of Σ_t . Recall that $n/p = \gamma$, which implies $O\left(\frac{1}{n}\right) = O\left(\frac{1}{p}\right)$. Then, it holds that

$$\begin{aligned}
T_2 &= \frac{1}{n} \mu_t^\top \hat{\Sigma}^+ \mu_t \\
&= \frac{\mu_t^\top}{\sqrt{n}} (\tilde{X}_n^\top \tilde{X}_n)^+ \frac{\mu_t}{\sqrt{n}} \\
&= \frac{\mu_t^\top}{\sqrt{n}} \sum_{i=1}^p \frac{1}{\sigma_i(\tilde{X}_n)^2} v_i(\tilde{X}_n) v_i(\tilde{X}_n)^\top \frac{\mu_t}{\sqrt{n}} \\
&= \frac{1}{\sigma_1(\tilde{X}_n)^2} \frac{\langle v_1(\tilde{X}_n), \mu_t \rangle^2}{n} + \frac{1}{\sigma_2(\tilde{X}_n)^2} \frac{\langle v_2(\tilde{X}_n), \mu_t \rangle^2}{n} + \sum_{i=3}^p \frac{1}{\sigma_i(\tilde{X}_n)^2} \frac{\langle v_i(\tilde{X}_n), \mu_t \rangle^2}{n} \\
&\leq \Theta\left(\frac{1}{p}\right) \left(1 - O\left(\frac{1}{p}\right)\right) + \frac{1}{c \cdot \tau} O\left(\frac{1}{p}\right) \\
&= O\left(\frac{1}{p}\right),
\end{aligned} \quad (\text{E.12})$$

where the penultimate inequality follows directly from (E.11) and Proposition E.2.

Finally, combining the bounds on the two terms we get

$$T_1 + T_2 = \frac{1}{n} \text{Tr}[\hat{\Sigma}_0^+ \Sigma_t] + O\left(\frac{1}{p}\right),$$

proving the claim. $\square$

We conclude this appendix with the proof of Theorem 3.1.

Proof of Theorem 3.1. Since $n > p$ in the considered regime, $\hat{\Sigma} = X^\top X/n$ is full rank almost surely, which implies that $\Pi = I - \hat{\Sigma}^+ \hat{\Sigma} = I - \hat{\Sigma}^{-1} \hat{\Sigma} = 0$. From (E.2), it follows that $B_X(\hat{\beta}; \beta) = 0$, so the risk is only characterized by the variance $V_X(\hat{\beta}; \beta)$. From this follows

$$R_X(\hat{\beta}, \beta) = V_X(\hat{\beta}; \beta) = \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}^+ (\Sigma_t + \mu_t \mu_t^\top)].$$

By directly applying Proposition E.3, it holds

$$\frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}^+ (\Sigma_t + \mu_t \mu_t^\top)] = \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}_0^+ \Sigma_t] + O\left(\frac{1}{p}\right),$$

where $\hat{\Sigma}_0 = \frac{X^{0\top} X^0}{\sqrt{n}}$. Plugging in the expression of $\frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}_0^+ \Sigma_t]$ given in [Yang et al., 2025, Theorem 3] gives the desired result. $\square$

E.3 Proof of Proposition C.1

Since we are in the setting where $n > p$, it holds that $B_X(\hat{\beta}; \beta) = 0$, which implies

$$R_X(\hat{\beta}, \beta) = V_X(\hat{\beta}; \beta) = \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}^+(\Sigma_t + \mu_t \mu_t^\top)].$$

Note that $\gamma_t = 0$ implies that $X_t = 0$ and $X = X_s$. We also note that (E.10) still holds for $X_t = 0$, implying that $\hat{\Sigma}$ is of rank p almost surely and, therefore, invertible. Thus, it holds

$$\text{Tr}[\hat{\Sigma}^+(\Sigma_t + \mu_t \mu_t^\top)] = \text{Tr}[\hat{\Sigma}^{-1}(\Sigma_t + \mu_t \mu_t^\top)].$$

To simplify exposition, we break this down into two terms

$$R_X(\hat{\beta}, \beta) = V_1 + V_2,$$

with $V_1 := \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}^{-1} \Sigma_t]$, $V_2 := \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}^{-1} \mu_t \mu_t^\top]$, and treat each of them separately.

Bounding the term V_2 . Note that $\gamma_t = 0$ implies $n = n_s$, so we will use these two values interchangeably throughout the proof. From the cyclic property of trace, we have

$$V_2 = \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}^{-1} \mu_t \mu_t^\top] = \sigma^2 \frac{\mu_t^\top \hat{\Sigma}^{-1} \mu_t}{n}.$$

Note that

$$\begin{aligned} \mu_t^\top \hat{\Sigma}^{-1} \mu_t &= \mu_t^\top \left(\frac{X^\top X}{n} \right)^{-1} \mu_t \\ &= \mu_t^\top \left(\frac{(Z_s \Sigma_s^{1/2} + 1_{n_s} \mu_s^\top)^\top (Z_s \Sigma_s^{1/2} + 1_{n_s} \mu_s^\top)}{n} \right)^{-1} \mu_t \\ &= \left(\Sigma_s^{-1/2} \mu_t \right)^\top \left(\frac{\left(Z_s + 1_{n_s} \left(\Sigma_s^{-1/2} \mu_s \right)^\top \right)^\top \left(Z_s + 1_{n_s} \left(\Sigma_s^{-1/2} \mu_s \right)^\top \right)}{n} \right)^{-1} \left(\Sigma_s^{-1/2} \mu_t \right) \\ &= \mu_t'^\top \hat{\Sigma}'^{-1} \mu_t', \end{aligned}$$

where we use the notation $\mu_t' := \Sigma_s^{-1/2} \mu_t$, $\mu_s' := \Sigma_s^{-1/2} \mu_s$ and $\hat{\Sigma}' := \frac{(Z_s + 1_{n_s} \mu_s'^\top)^\top (Z_s + 1_{n_s} \mu_s'^\top)}{n}$. Note that due to the assumed bound on the spectrum of Σ_s it holds that $\|\mu_t'\|_2 = O(\sqrt{p})$ and $\|\mu_s'\|_2 = O(\sqrt{p})$. Next, let us break down the vector μ_t' into its orthogonal projection onto the subspace $\{\mu_s'\}$ and $\{\mu_s'\}^\perp$ as

$$\mu_t' = \mu_{t\parallel s}' + \mu_{t\perp s}', \text{ where } \mu_{t\parallel s}' := \frac{\langle \mu_t', \mu_s' \rangle}{\|\mu_s'\|_2^2} \mu_s', \quad \mu_{t\perp s}' := \mu_t' - \mu_{t\parallel s}'. \quad (\text{E.13})$$

Moreover, as a decomposition into orthogonal spaces, it holds $\|\mu_{t\parallel s}'\|_2^2 + \|\mu_{t\perp s}'\|_2^2 = \|\mu_t'\|_2^2 = O(p)$. By using this decomposition, we will shift the focus from μ_t' to $\mu_{t\perp s}'$. Namely, it holds

$$\begin{aligned} V_2 &= \sigma^2 \frac{\mu_t'^\top \hat{\Sigma}'^{-1} \mu_t'}{n} = \sigma^2 \frac{(\mu_{t\parallel s}' + \mu_{t\perp s}')^\top \hat{\Sigma}'^{-1} (\mu_{t\parallel s}' + \mu_{t\perp s}')}{n} \\ &= \sigma^2 \frac{\mu_{t\perp s}'^\top \hat{\Sigma}'^{-1} \mu_{t\perp s}'}{\sqrt{n}} + 2\sigma^2 \frac{\mu_{t\perp s}'^\top \hat{\Sigma}'^{-1} \mu_{t\parallel s}'}{\sqrt{n}} + \frac{\mu_{t\parallel s}'^\top \hat{\Sigma}'^{-1} \mu_{t\parallel s}'}{\sqrt{n}} \\ &= \sigma^2 \frac{\mu_{t\perp s}'^\top \hat{\Sigma}'^{-1} \mu_{t\perp s}'}{\sqrt{n}} + O\left(\frac{1}{\sqrt{p}}\right), \end{aligned} \quad (\text{E.14})$$

where the last line follows from derivations analogous to the ones around (E.12), this time applying case 2. of Proposition E.2. To ease further exposition, we introduce $\tilde{\mu}_{t\perp s}' := \frac{\mu_{t\perp s}'}{\sqrt{n}}$, noting that $\|\tilde{\mu}_{t\perp s}'\|_2 = O(1)$.

In order to bound V_2 , we relate $\hat{\Sigma}'^{-1}$ to its zero-mean counterpart, as it is easier to work with mean zero data. Towards this end, we write out $\hat{\Sigma}'$ as

$$\begin{aligned}\hat{\Sigma}' &= \frac{(Z_s + 1_{n_s} \mu_s'^\top)^\top (Z_s + 1_{n_s} \mu_s'^\top)}{n} \\ &= \left(\frac{Z_s^\top Z_s}{n} + \frac{Z_s^\top 1_{n_s} \mu_s'^\top}{n} + \frac{\mu_s' 1_{n_s}^\top Z_s}{n} + \mu_s' \mu_s'^\top \right) \\ &= \left(\hat{\Sigma}'_0 + \frac{Z_s^\top 1_{n_s} \mu_s'^\top}{n} + \frac{\mu_s' 1_{n_s}^\top Z_s}{n} + \mu_s' \mu_s'^\top \right),\end{aligned}$$

for $\hat{\Sigma}'_0 := \frac{Z_s^\top Z_s}{n}$. All the terms above, except the first one, have rank 1, so we use Woodbury formula to take them out of the inverse when computing $\hat{\Sigma}'$. We introduce the following notation

$$\begin{aligned}u &:= \frac{\mu_s'}{\sqrt{n}}, & v &:= \frac{Z_s^\top 1_{n_s}}{\sqrt{n}}, \\ U &:= [u \ v] \in \mathbb{R}^{p \times 2}, \text{ and } C := \begin{bmatrix} n & 1 \\ 1 & 0 \end{bmatrix} \in \mathbb{R}^{2 \times 2}.\end{aligned}$$

Under this notation it holds

$$\frac{Z_s^\top 1_{n_s} \mu_s'^\top}{n} + \frac{\mu_s' 1_{n_s}^\top Z_s}{n} + \mu_s' \mu_s'^\top = U C U^\top.$$

Then, using Woodbury formula, we have

$$\begin{aligned}\hat{\Sigma}'^{-1} &= \left(\hat{\Sigma}'_0 + uv^\top + vu^\top + nuu^\top \right)^{-1} \\ &= \left(\hat{\Sigma}'_0 + U C U^\top \right)^{-1} \\ &= \hat{\Sigma}'_0^{-1} - \hat{\Sigma}'_0^{-1} U (C^{-1} - U^\top \hat{\Sigma}'_0^{-1} U)^{-1} U^\top \hat{\Sigma}'_0^{-1}.\end{aligned}$$

We now compute the 2×2 block

$$C^{-1} - U^\top \hat{\Sigma}'_0^{-1} U = \begin{bmatrix} -u^\top \hat{\Sigma}'_0^{-1} u & 1 - u^\top \hat{\Sigma}'_0^{-1} v \\ 1 - v^\top \hat{\Sigma}'_0^{-1} u & -n - v^\top \hat{\Sigma}'_0^{-1} v \end{bmatrix} = \begin{bmatrix} -a & 1 - b \\ 1 - b & -n - d \end{bmatrix},$$

where

$$a := u^\top \hat{\Sigma}'_0^{-1} u, \quad b := v^\top \hat{\Sigma}'_0^{-1} u = u^\top \hat{\Sigma}'_0^{-1} v, \quad d := v^\top \hat{\Sigma}'_0^{-1} v.$$

Hence

$$(C^{-1} - U^\top \hat{\Sigma}'_0^{-1} U)^{-1} = \frac{1}{\Delta} \begin{bmatrix} -n - d & b - 1 \\ b - 1 & -a \end{bmatrix}, \quad \Delta := a(n + d) - (1 - b)^2.$$

Plugging back and simplifying gives the explicit formula:

$$\hat{\Sigma}'^{-1} = \hat{\Sigma}'_0^{-1} - \frac{1}{\Delta} \hat{\Sigma}'_0^{-1} ((-n - d)uu^\top - (1 - b)(uv^\top + vu^\top) - a vv^\top) \hat{\Sigma}'_0^{-1},$$

which is valid whenever $\Delta \neq 0$, i.e., whenever $C^{-1} - U^\top \hat{\Sigma}'_0^{-1} U$ is invertible.

We will now analyze the a, b, d terms. First, for some constants $c_1, c_2 > 0$ it holds almost surely that

$$c_2 > \lambda_1(\hat{\Sigma}'_0) \geq \lambda_p(\hat{\Sigma}'_0) \geq c_1 > 0, \tag{E.15}$$

which follows from Bai–Yin theorem [Bai and Silverstein, 2010, Theorem 5.11], as Z_s has i.i.d entries with mean zero, unit variance and bounded fourth moments. From this, it follows that

$$\begin{aligned} |a| &= \left| u^\top \hat{\Sigma}'_0^{-1} u \right| \\ &= \left\| \frac{\mu'_s}{\sqrt{n}} \hat{\Sigma}'_0^{-1} \frac{\mu'_s}{\sqrt{n}} \right\|_2 \\ &\leq \left\| \frac{\mu'_s}{\sqrt{n}} \right\|_2 \left\| \hat{\Sigma}'_0^{-1} \right\|_2 \left\| \frac{\mu'_s}{\sqrt{n}} \right\|_2 \\ &\leq c, \end{aligned}$$

as well as

$$|a| \geq c \cdot (\lambda_1(\hat{\Sigma}'_0))^{-1} \geq c > 0,$$

Similarly, we have

$$|b| = \left| v^\top \hat{\Sigma}'_0^{-1} u \right| = \left\| \frac{\mu'_s}{\sqrt{n}} \hat{\Sigma}'_0^{-1} \frac{Z_s^\top 1_{n_s}}{\sqrt{n}} \right\|_2 \leq \left\| \frac{\mu'_s}{\sqrt{n}} \right\|_2 \left\| \hat{\Sigma}'_0^{-1} \right\|_2 \left\| \frac{Z_s^\top 1_{n_s}}{\sqrt{n}} \right\|_2 \leq c\sqrt{p}, \quad (\text{E.16})$$

where the last inequality follows with high probability over the sampling of Z_s , since $\frac{Z_s^\top 1_{n_s}}{\sqrt{n}}$ is a vector with p i.i.d entries of mean zero and $O(1)$ variance. Finally, we have

$$\begin{aligned} |d| &= \left| v^\top \hat{\Sigma}'_0^{-1} v \right| \\ &= \left\| \frac{1_{n_s}^\top Z_s}{\sqrt{n}} \hat{\Sigma}'_0^{-1} \frac{Z_s^\top 1_{n_s}}{\sqrt{n}} \right\|_2 \\ &\leq \left\| \frac{Z_s^\top 1_{n_s}}{\sqrt{n}} \right\|_2 \left\| \hat{\Sigma}'_0^{-1} \right\|_2 \left\| \frac{Z_s^\top 1_{n_s}}{\sqrt{n}} \right\|_2 \\ &\leq cp, \end{aligned}$$

again with high probability.

We can now prove that, with high probability, $\Delta = \Omega(p)$. Using Cauchy-Schwarz, it holds that

$$b^2 = |\langle u, v \rangle|_{A^{-1}} \leq \|u\|_{A^{-1}} \|v\|_{A^{-1}} = ad,$$

from which it follows that

$$\Delta = a(n + d) - (1 - b)^2 \geq an - 1 + 2b = \Omega(p), \quad (\text{E.17})$$

since a is lower bounded by a constant and $|b| \leq c\sqrt{p}$.

Turning back to the value of interest, we write out

$$\begin{aligned} \tilde{\mu}_{t \perp s}^\top \hat{\Sigma}'_0^{-1} \tilde{\mu}_{t \perp s} &= \tilde{\mu}_{t \perp s}^\top \hat{\Sigma}'_0^{-1} \tilde{\mu}_{t \perp s} \\ &\quad - \tilde{\mu}_{t \perp s}^\top \frac{1}{\Delta} \hat{\Sigma}'_0^{-1} ((-n - d)uu^\top + (1 - b)(uv^\top + vu^\top) - a vv^\top) \hat{\Sigma}'_0^{-1} \tilde{\mu}_{t \perp s} \\ &= \tilde{\mu}_{t \perp s}^\top \hat{\Sigma}'_0^{-1} \tilde{\mu}_{t \perp s} + T_{u,u} + T_{u,v} + T_{v,v}, \end{aligned}$$

where $T_{u,u}$ is the summand corresponding to $uu^\top$, $T_{u,v}$ to $uv^\top + vu^\top$, and $T_{v,v}$ to $vv^\top$. We will prove that each of these terms, except for $\tilde{\mu}_{t \perp s}^\top \hat{\Sigma}'_0^{-1} \tilde{\mu}_{t \perp s}$, is vanishing.

First, we state a useful claim, that for arbitrary deterministic unit vectors $w_1 \in \mathbb{R}^p$ and $w_2 \in \mathbb{R}^p$ it holds with overwhelming probability

$$w_1^\top \hat{\Sigma}'_0^{-1} w_2 = \frac{\gamma}{\gamma - 1} \langle w_1, w_2 \rangle + O(n^{-c_1}), \quad (\text{E.18})$$

for some constant $c_1 > 0$.

Proof of claim in (E.18). The result follows directly from [Yang et al., 2025, Theorem 27]. For clarity, we refer to the relevant parts of Section B.3.1 of that work. While Theorem 27 is stated in the more general anisotropic setting, it specializes to our isotropic case by taking Λ , U and V from (B.3) from their work to be the identity. Substituting these choices into equation (B.6) from Yang

et al. [2025] for $z = 0$, implies

$$\alpha_1(0) + \alpha_2(0) = 1 - \frac{p}{n} = \frac{\gamma - 1}{\gamma}.$$

Substituting this into (B.7) and applying Theorem 27 from the mentioned paper, yields with overwhelming probability

$$\left| w_1^\top \hat{\Sigma}'_0^{-1} w_2 - w_1^\top \frac{\gamma}{\gamma - 1} I_p w_2 \right| \leq n^{-c_1},$$

for any $c_1 < -1/2 + 2/\psi$. Recalling that Z_s has its ψ -th moment bounded for $\psi > 4$, implies $c_1 > 0$. ♣

We can now use (E.18) to tackle the terms $T_{u,u}$ and $T_{u,v}$. Namely, we have that

$$\begin{aligned} \tilde{\mu}_{t\perp s}^\top \hat{\Sigma}'_0^{-1} u &= \|\tilde{\mu}_{t\perp s}\|_2 \|u\|_2 \left(\frac{\gamma}{\gamma - 1} \langle \tilde{\mu}_{t\perp s}, u \rangle + O(n^{-c_1}) \right) \\ &= c \left(\frac{\gamma}{\gamma - 1} \left\langle \tilde{\mu}_{t\perp s}, \frac{\mu'_s}{\sqrt{n}} \right\rangle + O(n^{-c_1}) \right) \\ &= O(n^{-c_1}), \end{aligned}$$

with high probability. From this, it follows that

$$T_{u,u} = \frac{n+d}{\Delta} \tilde{\mu}_{t\perp s}^\top \hat{\Sigma}'_0^{-1} u u^\top \hat{\Sigma}'_0^{-1} \tilde{\mu}_{t\perp s} = O(n^{-2c_1}). \quad (\text{E.19})$$

Similarly,

$$\begin{aligned} |T_{u,v}| &= \left| \frac{1-b}{\Delta} \tilde{\mu}_{t\perp s}^\top \hat{\Sigma}'_0^{-1} (uv^\top + vu^\top) \hat{\Sigma}'_0^{-1} \tilde{\mu}_{t\perp s} \right| \\ &= \left| 2(\tilde{\mu}_{t\perp s}^\top \hat{\Sigma}'_0^{-1} u) \cdot \left(\frac{1-b}{\Delta} v^\top \hat{\Sigma}'_0^{-1} \tilde{\mu}_{t\perp s} \right) \right| \\ &\leq O(n^{-c_1}) \cdot \left\| \frac{1-b}{\Delta} \frac{Z_s^\top 1_{n_s}}{\sqrt{n}} \right\|_2 \left\| \hat{\Sigma}'_0^{-1} \right\|_2 \|\tilde{\mu}_{t\perp s}\|_2 \\ &= O(n^{-c_1}), \end{aligned} \quad (\text{E.20})$$

where the last inequality holds with high probability due to (E.15), (E.16), and (E.17).

Let us denote by $\tilde{1}_{n_s} := \frac{1_{n_s}}{\sqrt{n}}$ and turn to the term $T_{v,v}$.

Notice that

$$\begin{aligned} T_{v,v} &= \frac{a}{\Delta} \tilde{\mu}_{t\perp s}^\top \hat{\Sigma}'_0^{-1} v v^\top \hat{\Sigma}'_0^{-1} \tilde{\mu}_{t\perp s} \\ &= \frac{n a}{\Delta} \left(\tilde{1}_{n_s}^\top \frac{Z_s}{\sqrt{n}} \left(\frac{Z_s^\top Z_s}{n} \right)^{-1} \tilde{\mu}_{t\perp s} \right)^2 \\ &= c \left(\tilde{1}_{n_s}^\top \frac{Z_s}{\sqrt{n}} \left(\frac{Z_s^\top Z_s}{n} \right)^{-1} \tilde{\mu}_{t\perp s} \right)^2. \end{aligned} \quad (\text{E.21})$$

Let us introduce a matrix $Q = [q_1 \ \dots \ q_p] \in \mathbb{R}^{p \times p}$, whose columns form an orthonormal basis, such that $q_1 = \frac{\tilde{\mu}_{t\perp s}}{\|\tilde{\mu}_{t\perp s}\|_2}$. Then, we have that

$$\begin{aligned} \tilde{1}_{n_s}^\top \frac{Z_s}{\sqrt{n}} \left(\frac{Z_s^\top Z_s}{n} \right)^{-1} \tilde{\mu}_{t\perp s} &= \tilde{1}_{n_s}^\top \frac{Z_s}{\sqrt{n}} Q Q^\top \left(\frac{Z_s^\top Z_s}{n} \right)^{-1} \tilde{\mu}_{t\perp s} \\ &= \sum_{k=1}^p \tilde{1}_{n_s}^\top \frac{Z_s}{\sqrt{n}} q_k \cdot q_k^\top \left(\frac{Z_s^\top Z_s}{n} \right)^{-1} \tilde{\mu}_{t\perp s}. \end{aligned} \quad (\text{E.22})$$

Using (E.18) and a union bound, it holds with overwhelming probability that

$$q_k \left(\frac{Z_s^\top Z_s}{n} \right)^{-1} \tilde{\mu}_{t\perp s} = \frac{\gamma}{\gamma-1} \langle q_k, \tilde{\mu}_{t\perp s} \rangle + O(n^{-c_1}) = \begin{cases} \frac{\gamma}{\gamma-1} \|\tilde{\mu}_{t\perp s}\|_2 + O(n^{-c_1}), & k = 1, \\ O(n^{-c_1}), & k > 1. \end{cases}$$

Plugging this into (E.22) yields

$$\tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} \left(\frac{Z_s^\top Z_s}{n} \right)^{-1} \tilde{\mu}_{t\perp s} = \tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} \tilde{\mu}_{t\perp s} \cdot \frac{\gamma}{\gamma-1} + O(n^{-c_1}) \cdot \sum_{k=1}^p \tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} q_k. \quad (\text{E.23})$$

Let us first analyze the mean and variance of the random variable $\tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} \tilde{\mu}_{t\perp s}$, namely,

$$\begin{aligned} \mathbb{E} \left[\tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} \tilde{\mu}_{t\perp s} \right] &= \mathbb{E} \left[\frac{1}{\sqrt{n}} \sum_{i=1}^n \sum_{j=1}^p Z_{i,j} (\tilde{\mathbf{1}}_{n_s})_i (\tilde{\mu}_{t\perp s})_j \right] = 0, \\ \text{Var} \left(\tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} \tilde{\mu}_{t\perp s} \right) &= \text{Var} \left(\frac{1}{\sqrt{n}} \sum_{i=1}^n \sum_{j=1}^p Z_{i,j} (\tilde{\mathbf{1}}_{n_s})_i (\tilde{\mu}_{t\perp s})_j \right) \\ &= \frac{1}{n} \|\tilde{\mathbf{1}}_{n_s}\|_2^2 \|\tilde{\mu}_{t\perp s}\|_2^2 = O\left(\frac{1}{n}\right). \end{aligned}$$

Therefore, using Chebyshev inequality, we have that

$$\left| \tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} \tilde{\mu}_{t\perp s} \right| = O(n^{-c_2}),$$

with high probability, for some constant $1/2 > c_2 > 0$. Similarly, we calculate the mean and variance of the random variable $\sum_{k=1}^p \tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} q_k$ as

$$\begin{aligned} \mathbb{E} \left[\sum_{k=1}^p \tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} q_k \right] &= \mathbb{E} \left[\frac{1}{\sqrt{n}} \sum_{k=1}^p \sum_{i=1}^n \sum_{j=1}^p Z_{i,j} (\tilde{\mathbf{1}}_{n_s})_i (q_k)_j \right] = 0, \\ \text{Var} \left(\sum_{k=1}^p \tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} q_k \right) &= \text{Var} \left(\frac{1}{\sqrt{n}} \sum_{k=1}^p \sum_{i=1}^n \sum_{j=1}^p Z_{i,j} (\tilde{\mathbf{1}}_{n_s})_i (q_k)_j \right) \\ &= \frac{1}{n} \sum_{k=1}^p \|\tilde{\mathbf{1}}_{n_s}\|_2^2 \|q_k\|_2^2 = O(1). \end{aligned}$$

Again, Chebyshev inequality implies

$$\left| O(n^{-c_1}) \cdot \sum_{k=1}^p \tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} q_k \right| = O(n^{-c_1/2}),$$

with high probability. Plugging the obtained results into (E.23) and using a union bound on the probabilities, we get that

$$\begin{aligned} \left| \tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} \left(\frac{Z_s^\top Z_s}{n} \right)^{-1} \tilde{\mu}_{t\perp s} \right| &\leq \left| O(n^{-c_1}) \cdot \sum_{k=1}^p \tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} q_k \right| + \left| \tilde{\mathbf{1}}_{n_s}^\top \frac{Z_s}{\sqrt{n}} \tilde{\mu}_{t\perp s} \right| \\ &= O(n^{-c_1/2}), \end{aligned}$$

with high probability. Then, we directly obtain a bound for (E.21) in the form of

$$T_{v,v} = O(n^{-c_1}), \quad (\text{E.24})$$

which holds for some constant $c_1 > 0$ with high probability. Combining the bound in (E.18) and the three bounds on the terms (E.19), (E.20) and (E.24), we get

$$\tilde{\mu}_{t\perp s}^\top \hat{\Sigma}'^{-1} \tilde{\mu}_{t\perp s} = \frac{\gamma}{\gamma-1} \|\tilde{\mu}_{t\perp s}\|_2^2 + O(n^{-c}), \quad (\text{E.25})$$

for some $c > 0$. Using this in (E.14) yields

$$V_2 = \sigma^2 \frac{\gamma}{\gamma - 1} \|\tilde{\mu}_{t \perp s}\|_2^2 + O(n^{-c}),$$

with high probability. Lastly, note that

$$\|\tilde{\mu}_{t \perp s}\|_2^2 = \frac{1}{n} \|\mu'_{t \perp s}\|_2^2 = \frac{1}{n} \left(\|\mu'_t\|_2^2 - \|\mu'_t\|_s^2 \right) = \frac{1}{n} \left(\|\Sigma_s^{-1/2} \mu_t\|_2^2 - \left(\frac{\mu_t^\top \Sigma_s^{-1} \mu_s}{\|\Sigma_s^{-1/2} \mu_s\|_2} \right)^2 \right).$$

Bounding the term V_1 . By following exactly the proof of the bound of the term T_1 in Proposition E.3, one directly gets the same conclusion that

$$V_1 = \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}_0^{-1} \Sigma_t] + O\left(\frac{1}{p}\right).$$

Notice that

$$\hat{\Sigma}_0 = \frac{X^\top X}{n} = \frac{X_s^\top X_s}{n} = \frac{\Sigma_s^{1/2} Z_s^\top Z_s \Sigma_s^{1/2}}{n}.$$

Thus,

$$\frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}_0^{-1} \Sigma_t] = \frac{\sigma^2}{n} \text{Tr} \left[\Sigma_s^{-1/2} \left(\frac{Z_s^\top Z_s}{n} \right)^{-1} \Sigma_s^{-1/2} \Sigma_t \right] = \frac{\sigma^2}{n} \text{Tr} \left[\Sigma_s^{-1/2} \Sigma_t \Sigma_s^{-1/2} \left(\frac{Z_s^\top Z_s}{n} \right)^{-1} \right].$$

Let us write the SVD of $\Sigma_s^{-1/2} \Sigma_t \Sigma_s^{-1/2}$ as

$$\Sigma_s^{-1/2} \Sigma_t \Sigma_s^{-1/2} = \sum_{i=1}^p \lambda_i (\Sigma_s^{-1/2} \Sigma_t \Sigma_s^{-1/2}) w_i w_i^\top,$$

where $w_i := v_i(\Sigma_s^{-1/2} \Sigma_t \Sigma_s^{-1/2})$. Then it holds with overwhelming probability

$$\begin{aligned} \frac{\sigma^2}{n} \text{Tr} \left[\Sigma_s^{-1/2} \Sigma_t \Sigma_s^{-1/2} \left(\frac{Z_s^\top Z_s}{n} \right)^{-1} \right] &= \frac{\sigma^2}{n} \sum_{i=1}^p \lambda_i (\Sigma_s^{-1/2} \Sigma_t \Sigma_s^{-1/2}) w_i^\top \left(\frac{Z_s^\top Z_s}{n} \right)^{-1} w_i \\ &= \frac{\sigma^2}{n} \sum_{i=1}^p \lambda_i (\Sigma_s^{-1/2} \Sigma_t \Sigma_s^{-1/2}) \frac{\gamma}{\gamma - 1} \left(\|w_i\|_2^2 + O(n^{-c}) \right) \\ &= \left(\frac{\sigma^2}{n} \frac{\gamma}{\gamma - 1} \sum_{i=1}^p \lambda_i (\Sigma_s^{-1/2} \Sigma_t \Sigma_s^{-1/2}) \right) + O(n^{-c}) \\ &= \frac{\sigma^2}{n} \frac{\gamma}{\gamma - 1} \text{Tr}(\Sigma_t \Sigma_s^{-1}) + O(n^{-c}), \end{aligned}$$

where the second line holds with overwhelming probability by using (E.18) and the union bound. The previous bound also holds with high probability, since overwhelming probability implies it.

Finally, by combining the bounds on V_1 and V_2 , one gets that, with high probability,

$$\left| R_X(\hat{\beta}, \beta) - \frac{\sigma^2}{n} \frac{\gamma}{\gamma - 1} \text{Tr}(\Sigma_t \Sigma_s^{-1}) - \frac{\sigma^2}{n} \frac{\gamma}{\gamma - 1} \left(\|\Sigma_s^{-1/2} \mu_t\|_2^2 - \left(\frac{\mu_t^\top \Sigma_s^{-1} \mu_s}{\|\Sigma_s^{-1/2} \mu_s\|_2} \right)^2 \right) \right| = O(n^{-c}),$$

for some constant $c > 0$. Taking the limit $n \rightarrow \infty$ on both sides yields the desired result.

E.4 Proof of Theorem 3.2

We start by removing α_2 from the fixed point in (3.2) and replacing it by $1 - \frac{p}{n} - \alpha_1$. We rename α_1 as α for convenience. Plugging this into the definition of $\mathcal{R}_u(M)$, we get

$$\mathcal{R}_u(M) = \frac{\sigma^2}{n} \text{Tr} \left[(\alpha_1 M + \alpha_2 \text{Id}_{p \times p})^{-1} \right] = \frac{\sigma^2}{n} \sum_{i=1}^p \frac{1}{\lambda_i \alpha + 1 - \frac{p}{n} - \alpha},$$

where as in Theorem 3.1 we refer to $\lambda_1 \geq \dots \geq \lambda_p$ as the eigenvalues of the matrix M . Furthermore, the fixed point equation (3.2) can be rewritten as follows:

$$\sum_{i=1}^p \frac{1}{\lambda_i \alpha + 1 - \frac{p}{n} - \alpha} = \frac{p + n\alpha - n_s}{1 - \frac{p}{n} - \alpha} = n \left(\frac{n - n_s}{n - p - n\alpha} - 1 \right). \quad (\text{E.26})$$

Thus, we have

$$\mathcal{R}_u(M) = \frac{\sigma^2}{n} \cdot n \left(\frac{n - n_s}{n - p - n\alpha} - 1 \right) = \sigma^2 \left(\frac{1 - \frac{n_s}{n}}{1 - \frac{p}{n} - \alpha} - 1 \right). \quad (\text{E.27})$$

Now, due to the RHS of (E.27), it can be seen that $\mathcal{R}_u(M)$ is an increasing function of α . Let us denote by $\vec{\lambda} := [\lambda_1, \dots, \lambda_p]$. Then, for fixed n, p, n_s and $\vec{\lambda}$, we will refer to $\alpha(\vec{\lambda})$ as the solution to the fixed point equation (E.26). Note that following Yang et al. [2025][Appendix B.3.2] we have that this solutions is unique and $0 < \alpha(\vec{\lambda}) < \frac{n-p}{n}$.

Consider a function $f : \mathbb{R}_{\geq 0}^p \rightarrow \mathbb{R}_{\geq 0}^p$. We call a function f *good*, if and only if

$$\sum_{i=1}^p \frac{1}{f(\vec{\lambda})_i \alpha(\vec{\lambda}) + 1 - \frac{p}{n} - \alpha(\vec{\lambda})} < \sum_{i=1}^p \frac{1}{\lambda_i \alpha(\vec{\lambda}) + 1 - \frac{p}{n} - \alpha(\vec{\lambda})}. \quad (\text{E.28})$$

We claim that if f is good, then

$$\alpha(f(\vec{\lambda})) < \alpha(\vec{\lambda}). \quad (\text{E.29})$$

Proof of the claim. Consider a good function f . Then, we have

$$\sum_{i=1}^p \frac{1}{f(\vec{\lambda})_i \alpha(\vec{\lambda}) + 1 - \frac{p}{n} - \alpha(\vec{\lambda})} < \sum_{i=1}^p \frac{1}{\lambda_i \alpha(\vec{\lambda}) + 1 - \frac{p}{n} - \alpha(\vec{\lambda})} = n \left(\frac{n - n_s}{n - p - n\alpha(\vec{\lambda})} - 1 \right).$$

Furthermore, setting $\alpha = 0$ we get

$$\begin{aligned} \sum_{i=1}^p \frac{1}{f(\vec{\lambda})_i \cdot 0 + 1 - \frac{p}{n} - 0} &= p \frac{n}{n - p} \\ &> n \frac{p - n_s}{n - p} \\ &= n \left(\frac{n - n_s}{n - p - n \cdot 0} - 1 \right). \end{aligned}$$

By continuity, there exists $\alpha_0 \in (0, \alpha(\vec{\lambda}))$ for which

$$\sum_{i=1}^p \frac{1}{f(\vec{\lambda})_i \alpha_0 + 1 - \frac{p}{n} - \alpha_0} = n \left(\frac{n - n_s}{n - p - n\alpha_0} - 1 \right),$$

implying $\alpha(f(\vec{\lambda})) = \alpha_0 < \alpha(\vec{\lambda})$, which concludes the proof. ♣

Next, for $i, j \in [p]$ s.t. $i < j$, we introduce a function $f_c^{i,j} : \mathbb{R}_{\geq 0}^p \rightarrow \mathbb{R}_{\geq 0}^p$ defined as

$$f_c^{i,j}(\vec{\lambda})_k = \begin{cases} \lambda_i - c & k = i, \\ \lambda_j + c & k = j, \\ \lambda_k & k \neq i, j, \end{cases}$$

where $c > 0$. We now claim that $f_c^{i,j}$ is good for any $i, j \in [p]$ and $c > 0$, such that $\lambda_i > \lambda_j + c$.

Proof of the claim. The claim is equivalent to

$$\begin{aligned} &\frac{1}{(\lambda_i - c)\alpha(\vec{\lambda}) + 1 - \frac{p}{n} - \alpha(\vec{\lambda})} + \frac{1}{(\lambda_j + c)\alpha(\vec{\lambda}) + 1 - \frac{p}{n} - \alpha(\vec{\lambda})} \\ &< \frac{1}{\lambda_i \alpha(\vec{\lambda}) + 1 - \frac{p}{n} - \alpha(\vec{\lambda})} + \frac{1}{\lambda_j \alpha(\vec{\lambda}) + 1 - \frac{p}{n} - \alpha(\vec{\lambda})}. \end{aligned}$$

For simplicity, let $\delta := 1 - \frac{p}{n} - \alpha(\vec{\lambda})$ and $\alpha := \alpha(\vec{\lambda})$. Then,

$$\begin{aligned}
& \frac{1}{(\lambda_i - c)\alpha + \delta} + \frac{1}{(\lambda_j + c)\alpha + \delta} < \frac{1}{\lambda_i\alpha + \delta} + \frac{1}{\lambda_j\alpha + \delta} \\
& \iff \frac{\alpha(\lambda_i + \lambda_j) + 2\delta}{(\lambda_i\alpha - c\alpha + \delta)(\lambda_j\alpha + c\alpha + \delta)} < \frac{\alpha(\lambda_i + \lambda_j) + 2\delta}{(\lambda_i\alpha + \delta)(\lambda_j\alpha + \delta)} \\
& \iff (\lambda_i\alpha + \delta)(\lambda_j\alpha + \delta) < (\lambda_i\alpha - c\alpha + \delta)(\lambda_j\alpha + c\alpha + \delta) \\
& \iff c\alpha(\lambda_i\alpha + \delta) - c\alpha(\lambda_j\alpha + \delta) - c^2\alpha^2 > 0 \\
& \iff c\alpha^2(\lambda_i - \lambda_j) > c^2\alpha^2 \\
& \iff \lambda_i > \lambda_j + c,
\end{aligned}$$

which proves the claim. ♣

This implies that, for $t \in (0, 1)$, transformations of the form

$$(\lambda_i, \lambda_j) \rightarrow (t\lambda_i + (1-t)\lambda_j, (1-t)\lambda_i + t\lambda_j), \quad (\text{E.30})$$

are good.

Let us denote by $\vec{\lambda}' := [1, \dots, 1]$, which corresponds to eigenvalues of $M' := I_p \in \mathcal{M}$. Pick any $\vec{\lambda}'' \neq \vec{\lambda}'$ that corresponds to some matrix $M'' \in \mathcal{M}$, so it satisfies $\lambda_1'' \geq \lambda_2'' \geq \dots \geq \lambda_p''$, as well as $\sum_{i=1}^p \lambda_i'' = p$.

We recall the definition of *majorization*, as it will be used to conclude the proof. Namely, we say that $\vec{x} \in \mathbb{R}^p$ is *majorized* by $\vec{y} \in \mathbb{R}^p$ whenever for all $k \in [p]$

$$\sum_{i=1}^k x_i \leq \sum_{i=1}^k y_i,$$

and

$$\sum_{i=1}^p x_i = \sum_{i=1}^p y_i.$$

Firstly, we claim that $\vec{\lambda}'$ is majorized by $\vec{\lambda}''$. Suppose otherwise, that for some $k \in [p]$

$$\sum_{i=1}^k \lambda_i'' < \sum_{i=1}^k 1 = k,$$

implying also that $\lambda_k'' < 1$. Then, we have

$$p = \sum_{i=1}^p \lambda_i'' < (p-k)\lambda_k'' + k < (p-k) + k = p,$$

which is a contradiction.

Next, as $\vec{\lambda}'$ is majorized by $\vec{\lambda}''$, $\vec{\lambda}'$ can be derived from $\vec{\lambda}''$ by a finite sequence of steps of the form in (E.30) with $t \in [0, 1]$, see [Marshall et al., 1979, Chapter 4, Proposition A.1]. Since both vectors $\vec{\lambda}'$ and $\vec{\lambda}''$ are non-increasing, the $t = 0$ transformation can always be omitted. Moreover, $t = 1$ is just the identity transformation, so it can also be omitted and we actually have $t \in (0, 1)$. In formulas, we have that

$$\vec{\lambda}' = f_{c_l}^{i_l, j_l}(\dots f_{c_1}^{i_1, j_1}(\vec{\lambda}'') \dots).$$

Since each of the functions above is good, we have that $\alpha(\vec{\lambda}') < \alpha(\vec{\lambda}'')$. As $\mathcal{R}_u(M)$ is increasing with α , the smallest $\mathcal{R}_u(M)$ is achieved for $\vec{\lambda}' := [1, \dots, 1]$, that is, $I_p = \operatorname{arginf}_{M \in \mathcal{M}} \mathcal{R}_u(M)$.

E.5 Proof of $\mathcal{R}_u(\eta M) \leq \mathcal{R}_u(M)$

Consider the function $g_\eta : \mathbb{R}_{\geq 0}^p \rightarrow \mathbb{R}_{\geq 0}^p$ defined as $g_\eta(\vec{\lambda}) = \eta\vec{\lambda}$, for some $\eta > 1$. Note that, for all i ,

$$\frac{1}{g_\eta(\vec{\lambda})_i\alpha + 1 - \frac{p}{n} - \alpha} = \frac{1}{\eta\lambda_i\alpha + 1 - \frac{p}{n} - \alpha} < \frac{1}{\lambda_i\alpha + 1 - \frac{p}{n} - \alpha}.$$

Thus, $g_\eta(\vec{\lambda}) = \eta\vec{\lambda}$ is *good* in the sense of (E.28). From (E.29), we obtain that $\alpha(\eta\vec{\lambda}) < \alpha(\vec{\lambda})$. This implies the desired result as $\mathcal{R}_u$ is monotonically increasing in α from (E.27).

E.6 Coefficient defining system of equations of Theorem D.1

The (a_1, a_2, a_3, a_4) is the unique solution, with a_1, a_2 positive, to the following system of equations:

$$0 = 1 - \frac{1}{\gamma} \int \frac{a_1 \lambda^s + a_2 \lambda^t}{a_1 \lambda^s + a_2 \lambda^t + 1} d\hat{H}_p(\lambda^s, \lambda^t), \quad 0 = \frac{\gamma_s}{\gamma} - \frac{1}{\gamma} \int \frac{a_1 \lambda^s}{a_1 \lambda^s + a_2 \lambda^t + 1} d\hat{H}_p(\lambda^s, \lambda^t), \quad (\text{E.31})$$

$$a_1 + a_2 = -\frac{1}{\gamma} \int \frac{a_3 \lambda^s + a_4 \lambda^t}{(a_1 \lambda^s + a_2 \lambda^t + 1)^2} d\hat{H}_p(\lambda^s, \lambda^t), \quad a_1 = -\frac{1}{\gamma} \int \frac{a_3 \lambda^s + \lambda^s \lambda^t (a_3 a_2 - a_4 a_1)}{(a_1 \lambda^s + a_2 \lambda^t + 1)^2} d\hat{H}_p(\lambda^s, \lambda^t),$$

and (b_1, b_2, b_3, b_4) is the unique solution, with b_1, b_2 positive, to the following system of equations:

$$0 = 1 - \frac{1}{\gamma} \int \frac{b_1 \lambda^s + b_2 \lambda^t}{b_1 \lambda^s + b_2 \lambda^t + 1} d\hat{H}_p(\lambda^s, \lambda^t), \quad 0 = \frac{\gamma_s}{\gamma} - \frac{1}{\gamma} \int \frac{b_1 \lambda^s}{b_1 \lambda^s + b_2 \lambda^t + 1} d\hat{H}_p(\lambda^s, \lambda^t), \quad (\text{E.32})$$

$$0 = \int \frac{\lambda^s (b_3 - b_1 \lambda^t) + \lambda^t (b_4 - b_2 \lambda^t)}{(b_1 \lambda^s + b_2 \lambda^t + 1)^2} d\hat{H}_p(\lambda^s, \lambda^t), \quad 0 = \int \frac{\lambda^s (b_3 - b_1 \lambda^t) + \lambda^s \lambda^t (b_3 b_2 - b_4 b_1)}{(b_1 \lambda^s + b_2 \lambda^t + 1)^2} d\hat{H}_p(\lambda^s, \lambda^t).$$

E.7 Proof of Theorem D.1

Recall from (E.2) that bias and variance for non-zero centered data can be expressed as

$$B_X(\hat{\beta}; \beta) = \beta^\top \Pi (\Sigma_t + \mu_t \mu_t^\top) \Pi \beta \quad \text{and} \quad V_X(\hat{\beta}; \beta) = \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}^+ (\Sigma_t + \mu_t \mu_t^\top)],$$

where $\hat{\Sigma} = X^\top X/n$ and $\Pi = I - \hat{\Sigma}^+ \hat{\Sigma}$ (projection on the null space of X). To obtain the wanted result, we make a connection to zero-mean data and then use results from Song et al. [2024] to handle the zero-mean case. Unlike in the under-parametrized case, the bias term does not necessarily vanish. Thus, we start off by breaking it down into two terms

$$B_X(\hat{\beta}; \beta) = B_X^1(\hat{\beta}; \beta) + B_X^2(\hat{\beta}; \beta),$$

where $B_X^1(\hat{\beta}; \beta) = \beta^\top \Pi \Sigma_t \Pi \beta$ and $B_X^2(\hat{\beta}; \beta) = \beta^\top \Pi \mu_t \mu_t^\top \Pi \beta$. Moreover, we split the variance term as

$$V_X(\hat{\beta}; \beta) = V_X^1(\hat{\beta}; \beta) + V_X^2(\hat{\beta}; \beta),$$

with $V_X^1(\hat{\beta}; \beta) = \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}^+ \Sigma_t]$ and $V_X^2(\hat{\beta}; \beta) = \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}^+ \mu_t \mu_t^\top]$. We will deal with each of these terms individually.

Bounding the term $B_X^2(\hat{\beta}, \beta)$. Recall that $\tilde{X}_n = \frac{X}{\sqrt{n}}$. Then, similarly to (E.12), we can write the SVD of $\hat{\Sigma}$ as

$$\hat{\Sigma} = \sum_{i=1}^k \sigma_i^2(\tilde{X}_n) v_i(\tilde{X}_n) v_i(\tilde{X}_n)^\top,$$

where $k \leq \min(n, p) = n$ is the number of non-zero singular values of $\tilde{X}_n$. As in (E.10), we can conclude that $k = n$. Therefore, we have

$$I - \hat{\Sigma}^+ \hat{\Sigma} = I - \sum_{i=1}^n v_i(\tilde{X}_n) v_i(\tilde{X}_n)^\top = \sum_{i=n+1}^p v_i(\tilde{X}_n) v_i(\tilde{X}_n)^\top.$$

By definition, it holds that $\Pi \mu_t = (I - \hat{\Sigma}^+ \hat{\Sigma}) \mu_t$, from which it follows

$$\Pi \mu_t = \sum_{i=n+1}^p v_i(\tilde{X}_n) \langle v_i(\tilde{X}_n), \mu_t \rangle.$$

Due to Proposition E.2, it holds almost surely that

$$\left| \frac{\langle v_1(\tilde{X}_n), \mu_s \rangle}{\|\mu_s\|_2} \right|^2 + \left| \frac{\langle v_2(\tilde{X}_n), \mu_s \rangle}{\|\mu_s\|_2} \right|^2 \geq 1 - \frac{1}{c \cdot p},$$

from which it follows

$$\|\Pi\mu_t\|_2^2 = \sum_{i=n+1}^p \left| \left\langle v_i(\tilde{X}_n), \mu_t \right\rangle \right|^2 \leq \frac{1}{c \cdot p} \|\mu_t\|_2^2 = c.$$

Since β sampled independently from a sphere of constant radius and $\Pi\mu_t$ is of bounded norm, it is standard result that $|\langle \beta, \Pi\mu_t \rangle|^2$ is sub-exponential and, using Bernstein inequality, we can get that

$$B_X^2(\hat{\beta}, \beta) = \|\beta^\top \Pi\mu_t\|_2^2 = |\langle \beta, \Pi\mu_t \rangle|^2 = O\left(\frac{1}{p}\right), \quad (\text{E.33})$$

with high probability over the sampling of β .

Bounding the term $B_X^1(\hat{\beta}, \beta)$. We first introduce an object coming from a bias term of a ridge regression estimator with coefficient λ :

$$B_X^1(\lambda) := \lambda^2 \beta^\top (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t (\hat{\Sigma} + \lambda I)^{-1} \beta, \quad (\text{E.34})$$

defined for any $\lambda > 0$. It is more convenient to work with $B_X^1(\lambda)$ than $B_X^1(\hat{\beta}, \beta)$ and, in addition, $B_X^1(\lambda)$ approximates well $B_X^1(\hat{\beta}, \beta)$ for small λ . We formalize the second claim as

$$\left| B_X^1(\hat{\beta}, \beta) - B_X^1(\lambda) \right| = O(\lambda) \quad (\text{E.35})$$

proved in the same manner as [Song et al., 2024, D.82]. For convenience we also carry out the proof here.

Proof of the claim in (E.35). Let us write the SVD $\hat{\Sigma} = U D U^\top$. Moreover, we denote by $1_{D=0}$ and $1_{D>0}$ the diagonal matrices such that

$$(1_{D=0})_{i,i} = \begin{cases} 0, & D_{i,i} \neq 0 \\ 1, & D_{i,i} = 0 \end{cases} \quad (1_{D>0})_{i,i} = \begin{cases} 1, & D_{i,i} \neq 0 \\ 0, & D_{i,i} = 0 \end{cases}$$

Then it holds that

$$\begin{aligned} B_X^1(\hat{\beta}; \beta) &= \beta^\top (I - \hat{\Sigma}^+ \hat{\Sigma}) \Sigma_t (I - \hat{\Sigma}^+ \hat{\Sigma}) \beta \\ &= \beta^\top U 1_{D=0} U^\top \Sigma_t U 1_{D=0} U^\top \beta \\ &= \beta^\top U 1_{D=0} A 1_{D=0} U^\top \beta \\ &= \|A^{1/2} 1_{D=0} U^\top \beta\|_2^2, \end{aligned}$$

where we set $A := U^\top \Sigma_t U$. Furthermore, we have

$$\begin{aligned} B_X^1(\lambda) &= \lambda^2 \beta^\top (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t (\hat{\Sigma} + \lambda I)^{-1} \beta \\ &= \lambda^2 \beta^\top U (D + \lambda I)^{-1} A (D + \lambda I)^{-1} U^\top \beta \\ &= \|A^{1/2} \lambda (D + \lambda I)^{-1} U^\top \beta\|_2^2. \end{aligned}$$

Therefore, we have

$$\begin{aligned} \left| \sqrt{B_X^1(\hat{\beta}; \beta)} - \sqrt{B_X^1(\lambda)} \right| &\leq \|A^{1/2} (1_{D=0} - \lambda (D + \lambda I)^{-1}) U^\top \beta\|_2 \\ &\leq c \|A\|_2^{1/2} \|\lambda (D + \lambda I)^{-1} 1_{D>0}\|_2 \\ &\leq c \frac{\lambda}{\sigma_n(\hat{\Sigma})} = O(\lambda), \end{aligned}$$

where the third inequality holds as $\|A\|_2 = \|\Sigma_t\|_2 = O(1)$ and the last inequality follows from Proposition E.1 in the same manner as (E.11). Notice that $B_X^1(\lambda), B_X^1(\hat{\beta}; \beta) = O(1)$, since $\|\beta\|_2, \|\Sigma_t\|_2 = O(1)$ and $\sigma_n(\hat{\Sigma}) > c$. This finally implies

$$\left| B_X^1(\hat{\beta}; \beta) - B_X^1(\lambda) \right| = O(\lambda),$$

proving the claim. ♣

The next step is to prove the claim that, for $1 > \lambda > p^{-0.49}$, it holds that

$$B_X^1(\lambda) = \lambda^2 \beta^\top (\hat{\Sigma}_0 + \lambda I)^{-1} \Sigma_t (\hat{\Sigma}_0 + \lambda I)^{-1} \beta + O\left(\frac{\lambda^{-2}}{p}\right). \quad (\text{E.36})$$

Proof of the claim in (E.36). Towards this end, we have

$$\begin{aligned}\hat{\Sigma} &= \frac{1}{n}(X^\top X) \\ &= \frac{1}{n}(X^0 + 1_{n_t}\mu_t^\top + 1_{n_s}\mu_s^\top)^\top (X^0 + 1_{n_t}\mu_t^\top + 1_{n_s}\mu_s^\top) \\ &= \left(\frac{X^{0\top}X^0}{n} + \frac{X^{0\top}1_{n_t}\mu_t^\top}{n} + \frac{X^{0\top}1_{n_s}\mu_s^\top}{n} + \frac{\mu_t1_{n_t}^\top X^0}{n} + \frac{\mu_s1_{n_s}^\top X^0}{n} + \frac{\gamma_t}{\gamma}\mu_t\mu_t^\top + \frac{\gamma_s}{\gamma}\mu_s\mu_s^\top \right),\end{aligned}$$

where abusing notation we write $1_{n_s} = [1, \dots, 1, 0, \dots, 0]^\top \in \mathbb{R}^{n \times 1}$ (n_s ones followed by n_t zeros) and $1_{n_t} = [0, \dots, 0, 1, \dots, 1]^\top \in \mathbb{R}^{n \times 1}$ (n_s zeros followed by n_t ones).

All the terms above, except the first one, have rank 1, so we use Woodbury formula to take them out of the inverse when computing $(\hat{\Sigma} + \lambda I)^{-1}$. We consider the case $\varphi \neq 1$, as the case $\varphi = 1$ is analogous (it is in fact easier as some steps can be omitted). We first focus on the term $(\hat{\Sigma} + \lambda I)^{-1}$ and demonstrate how to handle $\frac{X^{0\top}1_{n_t}\mu_t^\top}{n} + \frac{\mu_t1_{n_t}^\top X^0}{n} + \frac{\gamma_t}{\gamma}\mu_t\mu_t^\top$. For this purpose, we introduce the following notation

$$\begin{aligned}A &:= \hat{\Sigma} + \lambda I - \frac{X^{0\top}1_{n_t}\mu_t^\top}{n} - \frac{\mu_t1_{n_t}^\top X^0}{n} - \frac{\gamma_t}{\gamma}\mu_t\mu_t^\top, \\ u &:= \frac{\mu_t}{\sqrt{n}}, \quad v := \frac{X^{0\top}1_{n_t}}{\sqrt{n}}, \\ U &:= [u \ v] \in \mathbb{R}^{p \times 2}, \text{ and } C := \begin{bmatrix} n\frac{\gamma_t}{\gamma} & 1 \\ 1 & 0 \end{bmatrix} \in \mathbb{R}^{2 \times 2}.\end{aligned}\tag{E.37}$$

Under this notation it holds

$$\frac{X^{0\top}1_{n_t}\mu_t^\top}{n} + \frac{\mu_t1_{n_t}^\top X^0}{n} + \frac{\gamma_t}{\gamma}\mu_t\mu_t^\top = UCU^\top.$$

Then, using Woodbury formula, we have

$$\begin{aligned}(\hat{\Sigma} + \lambda I)^{-1} &= \left(A + uv^\top + vu^\top + n\frac{\gamma_t}{\gamma}uu^\top \right)^{-1} \\ &= (A + UCU^\top)^{-1} \\ &= A^{-1} - A^{-1}U(C^{-1} - U^\top A^{-1}U)^{-1}U^\top A^{-1}.\end{aligned}$$

We now compute the 2×2 block

$$C^{-1} - U^\top A^{-1}U = \begin{bmatrix} -u^\top A^{-1}u & 1 - u^\top A^{-1}v \\ 1 - v^\top A^{-1}u & -n\frac{\gamma_t}{\gamma} - v^\top A^{-1}v \end{bmatrix} = \begin{bmatrix} -a & 1 - b \\ 1 - b & -n\frac{\gamma_t}{\gamma} - d \end{bmatrix},$$

where

$$a := u^\top A^{-1}u, \quad b := v^\top A^{-1}u = u^\top A^{-1}v, \quad d := v^\top A^{-1}v.\tag{E.38}$$

Hence

$$(C^{-1} - U^\top A^{-1}U)^{-1} = \frac{1}{\Delta} \begin{bmatrix} -n\frac{\gamma_t}{\gamma} - d & b - 1 \\ b - 1 & -a \end{bmatrix}, \quad \Delta := a \left(n\frac{\gamma_t}{\gamma} + d \right) - (1 - b)^2.\tag{E.39}$$

Plugging back and simplifying gives the explicit formula:

$$(\hat{\Sigma} + \lambda I)^{-1} = A^{-1} - \frac{1}{\Delta} A^{-1} \left(\left(-n\frac{\gamma_t}{\gamma} - d \right) uu^\top - (1 - b)(uv^\top + vu^\top) - a vv^\top \right) A^{-1},$$

which is valid whenever $\Delta \neq 0$, i.e., whenever $C^{-1} - U^\top A^{-1}U$ is invertible.

We will now analyze the a, b, d terms. First, recall that

$$A = \frac{X^{0\top}X^0}{n} + \frac{X^{0\top}1_{n_s}\mu_s^\top}{n} + \frac{\mu_s1_{n_s}^\top X^0}{n} + \frac{\gamma_s}{\gamma}\mu_s\mu_s^\top + \lambda I = \hat{\Sigma}_s + \lambda I,$$

where $\hat{\Sigma}_s := \frac{(X^0 + 1_{n_s} \mu_s^\top)^\top (X^0 + 1_{n_s} \mu_s^\top)}{n}$. Thus, we have

$$\|A^{-1}\|_2 \leq \lambda^{-1}.$$

From this, it follows that

$$|a| = |u^\top A^{-1} u| = \left\| \frac{\mu_t^\top}{\sqrt{n}} A^{-1} \frac{\mu_t}{\sqrt{n}} \right\|_2 \leq \left\| \frac{\mu_t}{\sqrt{n}} \right\|_2 \|A^{-1}\|_2 \left\| \frac{\mu_t}{\sqrt{n}} \right\|_2 \leq c \lambda^{-1}.$$

Similarly, we have

$$\begin{aligned} |b| &= |v^\top A^{-1} u| \\ &= \left\| \frac{\mu_t^\top}{\sqrt{n}} A^{-1} \frac{X^{0\top} 1_{n_t}}{\sqrt{n}} \right\|_2 \\ &\leq \left\| \frac{\mu_t}{\sqrt{n}} \right\|_2 \|A^{-1}\|_2 \left\| \frac{X^{0\top} 1_{n_t}}{\sqrt{n}} \right\|_2 \\ &\leq c \lambda^{-1} \sqrt{p}, \end{aligned}$$

where the last inequality follows with high probability over the sampling of X^0 , since $\frac{X^{0\top} 1_{n_t}}{\sqrt{n}}$ is a vector with p i.i.d entries of mean zero and $O(1)$ variance. Finally, we have

$$\begin{aligned} |d| &= |v^\top A^{-1} v| \\ &= \left\| \frac{1_{n_t}^\top X^0}{\sqrt{n}} A^{-1} \frac{X^{0\top} 1_{n_t}}{\sqrt{n}} \right\|_2 \\ &\leq \left\| \frac{X^{0\top} 1_{n_t}}{\sqrt{n}} \right\|_2 \|A^{-1}\|_2 \left\| \frac{X^{0\top} 1_{n_t}}{\sqrt{n}} \right\|_2 \\ &\leq c \lambda^{-1} p, \end{aligned}$$

again with high probability.

From a slight adjustment of the second part of Proposition E.2, it holds for the top singular value

$$\sigma_1(A) = \sigma_1(\hat{\Sigma}_s) + \lambda = \left(\sigma_1 \left(\frac{X^0 + 1_{n_s} \mu_s^\top}{\sqrt{n}} \right) \right)^2 + \lambda = \Theta(p),$$

and for the corresponding right singular vector

$$\left| \left\langle v_1(A), \frac{\mu_s}{\|\mu_s\|_2} \right\rangle \right| = \left| \left\langle v_1(\hat{\Sigma}_s), \frac{\mu_s}{\|\mu_s\|_2} \right\rangle \right| = \left| \left\langle v_1 \left(\frac{X^0 + 1_{n_s} \mu_s^\top}{\sqrt{n}} \right), \frac{\mu_s}{\|\mu_s\|_2} \right\rangle \right| = \sqrt{1 - O\left(\frac{1}{p}\right)}.$$

Note that, for $\varphi < 1$, it holds that $\left| \left\langle \frac{\mu_s}{\|\mu_s\|_2}, \frac{\mu_t}{\|\mu_t\|_2} \right\rangle \right| = \varphi < 1$. Using the triangle inequality and Cauchy-Schwarz gives

$$\left| \left\langle v_1(A), \frac{\mu_t}{\|\mu_t\|_2} \right\rangle \right| \leq \left| \left\langle \frac{\mu_s}{\|\mu_s\|_2}, \frac{\mu_t}{\|\mu_t\|_2} \right\rangle \right| + \left\| v_1(A) - \frac{\mu_s}{\|\mu_s\|_2} \right\|_2 \left\| \frac{\mu_t}{\|\mu_t\|_2} \right\|_2 \leq \varphi + O\left(\frac{1}{p}\right).$$

Therefore, it holds that

$$\begin{aligned}
|a| &= |u^\top A^{-1} u| = \left\| \frac{\mu_t^\top}{\sqrt{n}} A^{-1} \frac{\mu_t}{\sqrt{n}} \right\|_2 \\
&= \sum_{i=1}^p \frac{1}{\sigma_i(A)} \left| \left\langle v_i(A), \frac{\mu_t}{\sqrt{n}} \right\rangle \right|^2 \\
&= c \cdot \sum_{i=1}^p \frac{1}{\sigma_i(A)} \left| \left\langle v_i(A), \frac{\mu_t}{\|\mu_t\|_2} \right\rangle \right|^2 \\
&\geq c \sum_{i=2}^p \frac{1}{\sigma_i(A)} \left| \left\langle v_i(A), \frac{\mu_t}{\|\mu_t\|_2} \right\rangle \right|^2 \\
&\geq c \frac{1}{\sigma_2(A)} \sum_{i=2}^p \left| \left\langle v_i(A), \frac{\mu_t}{\|\mu_t\|_2} \right\rangle \right|^2 \\
&\geq c \left(1 - \left(\varphi + O\left(\frac{1}{p}\right) \right)^2 \right) > 0,
\end{aligned}$$

since $\sigma_2(A) = \sigma_2(\hat{\Sigma}_s) + \lambda = O(1)$ due to the second part of Proposition E.2. Note that, for $\varphi = 1$, we do not need this argument, as the μ_s terms are taken out of the inverse as well. In that case, we take $A = \left(\frac{X^{0\top} X^0}{n} + \lambda I \right)$, which immediately gives $\sigma_1(A) < c$.

We can now prove that, with high probability, $\Delta = \Omega(p)$. Using Cauchy-Schwarz, it holds that

$$b^2 = |\langle u, v \rangle|_{A^{-1}} \leq \|u\|_{A^{-1}} \|v\|_{A^{-1}} = ad,$$

from which it follows that

$$\Delta = a \left(n \frac{\gamma_t}{\gamma} + d \right) - (1 - b)^2 \geq an \frac{\gamma_t}{\gamma} - 1 + 2b = \Omega(p),$$

since a is lower bounded by a constant and $|b| \leq c\lambda^{-1}\sqrt{p} \leq cp^{0.99}$.

At this point, we have all the necessary bounds and we work towards proving the claim. We first expand the bias term

$$\begin{aligned}
B_X^1(\lambda) &= \lambda^2 \beta^\top (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t (\hat{\Sigma} + \lambda I)^{-1} \beta \\
&= \lambda^2 \beta^\top (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t (A + UCU^\top)^{-1} \beta \\
&= \lambda^2 \beta^\top (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t (A^{-1} - A^{-1}U(C^{-1} - U^\top A^{-1}U)^{-1}U^\top A^{-1}) \beta \\
&= \lambda^2 \beta^\top (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t A^{-1} \beta + S,
\end{aligned}$$

where $S := -\lambda^2 \beta^\top (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t A^{-1} U (C^{-1} - U^\top A^{-1}U)^{-1} U^\top A^{-1} \beta$.

We now prove that S is small. To do so, we decompose

$$\begin{aligned}
S &= -\lambda^2 \beta^\top (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t A^{-1} U (C^{-1} - U^\top A^{-1}U)^{-1} U^\top A^{-1} \beta \\
&= \lambda^2 \beta^\top (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t \frac{1}{\Delta} A^{-1} \left(\left(n \frac{\gamma_t}{\gamma} + d \right) uu^\top + (1 - b)(uv^\top + vu^\top) + a vv^\top \right) A^{-1} \beta \\
&= T_{u,u} + T_{u,v} + T_{v,v},
\end{aligned}$$

where $T_{u,u}$ is the summand corresponding to $uu^\top$, $T_{u,v}$ to $uv^\top + vu^\top$, and $T_{v,v}$ to $vv^\top$. Zooming in on one of the terms, it holds that

$$\begin{aligned}
T_{u,u} &= \lambda^2 \beta^\top (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t \frac{(n\gamma_t/\gamma + d)}{\Delta} A^{-1} uu^\top A^{-1} \beta \\
&= \left\langle \beta, \lambda^2 (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t \frac{(n\gamma_t/\gamma + d)}{\Delta} A^{-1} u \right\rangle \langle u^\top A^{-1}, \beta \rangle.
\end{aligned}$$

Note that

$$\left\| \lambda^2 (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t \frac{(n\gamma_t/\gamma + d)}{\Delta} A^{-1} u \right\|_2 \leq \lambda^2 \left\| (\hat{\Sigma} + \lambda I)^{-1} \right\|_2 \|\Sigma_t\|_2 \frac{(n\gamma_t/\gamma + d)}{\Delta} \|A^{-1}\|_2 \|u\|_2 \leq c\lambda^{-1},$$

and $\|u^\top A^{-1}\|_2 \leq c\lambda^{-1}$. Using this, we get that, with high probability, it holds

$$|T_{u,u}| \leq c \frac{\lambda^{-2}}{p}.$$

This is similar to how we obtained (E.33), since β is sampled independently from a sphere of constant radius. With analogous passages, we have that

$$|T_{u,v}| \leq c \frac{\lambda^{-2}}{p}, \quad |T_{v,v}| \leq c \frac{\lambda^{-2}}{p}$$

holds with high probability over the sampling of β . Putting all together, we get

$$B_X^1(\lambda) = \lambda^2 \beta^\top (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t A^{-1} \beta + O\left(\frac{\lambda^{-2}}{p}\right).$$

Using the same argumentation applied now to $(\hat{\Sigma} + \lambda I)^{-1}$ in $\lambda^2 \beta^\top (\hat{\Sigma} + \lambda I)^{-1} \Sigma_t A^{-1} \beta$ gives

$$B_X^1(\lambda) = \lambda^2 \beta^\top A^{-1} \Sigma_t A^{-1} \beta + O\left(\frac{\lambda^{-2}}{p}\right).$$

Lastly, doing all of this again to take out the terms containing μ_s from A , i.e., by taking

$$\tilde{A} := A - \frac{X^{0\top} 1_{n_s} \mu_s^\top}{n} - \frac{\mu_s 1_{n_s}^\top X^0}{n} - \frac{\gamma_s}{\gamma} \mu_s \mu_s^\top = \hat{\Sigma}_0 + \lambda I,$$

we get

$$B_X^1(\lambda) = \lambda^2 \beta^\top \tilde{A}^{-1} \Sigma_t \tilde{A}^{-1} \beta + O\left(\frac{\lambda^{-2}}{p}\right),$$

proving the claim. ♣

From [Song et al., 2024, D.82], it follows that

$$\left| \beta^\top \Pi_0 \Sigma_t \Pi_0 \beta - \lambda^2 \beta^\top (\hat{\Sigma}_0 + \lambda I)^{-1} \Sigma_t (\hat{\Sigma}_0 + \lambda I)^{-1} \beta \right| = O(\lambda), \quad (\text{E.40})$$

where $\Pi_0 = I - \hat{\Sigma}_0^+ \hat{\Sigma}_0$. Thus, by combining (E.35), (E.36) and (E.40), we conclude that

$$\left| B_X^1(\hat{\beta}, \beta) - \beta^\top \Pi_0 \Sigma_t \Pi_0 \beta \right| = O(\lambda) + O\left(\frac{\lambda^{-2}}{p}\right) = O(p^{-1/3}), \quad (\text{E.41})$$

where the last step is obtained by taking $p = \lambda^{-1/3}$ (this also satisfies $1 > \lambda > p^{-0.49}$, which was required to obtain (E.36)). As $B_X(\hat{\beta}, \beta) = B_X^1(\hat{\beta}, \beta) + B_X^2(\hat{\beta}, \beta)$ and $B_X^2(\hat{\beta}, \beta) = O(1/p)$ with high probability by (E.33), we conclude that

$$\left| B_X(\hat{\beta}, \beta) - \beta^\top \Pi_0 \Sigma_t \Pi_0 \beta \right| = O(p^{-1/3}) \quad (\text{E.42})$$

holds with high probability over the sampling of β and X . Plugging in the expression of $\beta^\top \Pi_0 \Sigma_t \Pi_0 \beta$ given in [Song et al., 2024, Theorem 4.1] yields, with high probability,

$$B_X(\hat{\beta}, \beta) = \int \frac{b_3 \lambda^s + (b_4 + 1) \lambda^t}{(b_1 \lambda^s + b_2 \lambda^t + 1)^2} d\hat{G}_p(\lambda^s, \lambda^t) + O(p^{-c}),$$

where (b_1, b_2, b_3, b_4) is the unique solution, with b_1, b_2 positive, to (E.32). Taking the limit $p, n \rightarrow \infty$ gives the desired result for the bias term.

Bounding the term $V_X^2(\hat{\beta}, \beta)$. Notice that the term $V_X^2(\hat{\beta}, \beta)$ coincides with T_2 from Proposition E.3. Moreover, we can follow the proof of the bound on T_2 verbatim, only substituting p for n in appropriate places (as we are now in an over-parametrized setting) to get

$$V_X^2(\hat{\beta}, \beta) = \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}_0^+ \mu_t \mu_t^\top] = O\left(\frac{1}{p}\right). \quad (\text{E.43})$$

Bounding the term $V_X^1(\hat{\beta}, \beta)$. To make a connection with zero-centered data, we will first prove that, with high probability, it holds

$$V_X^1(\hat{\beta}, \beta) = \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}^+ \Sigma_t] = \frac{1}{n} \text{Tr}[\hat{\Sigma}_0^+ \Sigma_t] + O\left(\frac{1}{p^{1/7}}\right). \quad (\text{E.44})$$

Similarly to the computation for $B_X^1(\hat{\beta}, \beta)$, we introduce an object coming from a variance term of a ridge regression estimator with coefficient λ :

$$V_X^1(\lambda) := \frac{1}{n} \text{Tr}[(\hat{\Sigma} + \lambda I)^{-2} \hat{\Sigma} \Sigma_t],$$

defined for any $\lambda > 0$. It is more convenient to work with $V_X^1(\lambda)$ than $V_X^1(\hat{\beta}, \beta)$ and, in addition, $V_X^1(\lambda)$ approximates $V_X^1(\hat{\beta}, \beta)$ well for small λ . We formalize the second claim as

$$\left| V_X^1(\hat{\beta}, \beta) - V_X^1(\lambda) \right| = O(\lambda), \quad (\text{E.45})$$

proved in the same manner as [Song et al., 2024, D.78]. For convenience we also carry out the proof here.

Proof of claim in (E.45). Let us write the SVD $\hat{\Sigma} = U D U^\top$. Then it holds that

$$\begin{aligned} V_X^1(\hat{\beta}, \beta) &= \frac{1}{n} \text{Tr}(U D^+ U^\top \Sigma_t), \\ V_X^1(\lambda) &= \frac{1}{n} \text{Tr}[U (D + \lambda I)^{-2} D U^\top \Sigma_t]. \end{aligned}$$

Therefore, we have

$$\begin{aligned} \left| V_X^1(\hat{\beta}, \beta) - V_X^1(\lambda) \right| &= \frac{1}{n} \left| \text{Tr} \left[U^\top \Sigma_t U (D^+ - (D + \lambda I)^{-2} D) \right] \right| \\ &\leq \|U^\top \Sigma_t U\|_2 \frac{1}{n} \sum_{i=1}^n \left[\frac{1}{\lambda_i(D)} - \frac{\lambda_i(D)}{(\lambda_i(D) + \lambda)^2} \right] \\ &\leq \frac{1}{\tau} \frac{2\lambda}{\lambda_n(D)^2} \\ &= c \cdot \frac{\lambda}{\lambda_n(\hat{\Sigma})^2} = O(\lambda). \end{aligned}$$

Here, we used the inequality $x^{-1} - (x + \lambda)^{-2} x \leq 2\lambda/x^2$ and the fact that $\hat{\Sigma}$ has n non-zero singular values, each bounded below by a constant, which follows from (E.10). This completes the proof of the claim. ♣

Relying on the derivations in [Song et al., 2024, D.2] we have that

$$V_X^1(\lambda) = \frac{d}{d\lambda} \left(\frac{\lambda}{n} \text{Tr} \left(\Sigma_t (\hat{\Sigma} + \lambda I)^{-1} \right) \right).$$

Let us denote by

$$\tilde{V}_X^1(\lambda) := \frac{\lambda}{n} \text{Tr} \left(\Sigma_t (\hat{\Sigma} + \lambda I)^{-1} \right).$$

We claim that, for any $t > 0$, it holds

$$\left| V_X^1(\lambda) - \frac{1}{t\lambda} \left(\tilde{V}_X^1(\lambda + t\lambda) - \tilde{V}_X^1(\lambda) \right) \right| = O(t\lambda^{-2}). \quad (\text{E.46})$$

Proof of claim in E.46. We begin by transforming the LHS:

$$\begin{aligned}
\frac{1}{t\lambda}(\tilde{V}_X^1(\lambda + t\lambda) - \tilde{V}_X^1(\lambda)) &= \frac{1}{n} \text{Tr} \left(\Sigma_t \frac{1}{t\lambda} \left((\lambda + t\lambda) (\hat{\Sigma} + (\lambda + t\lambda)I)^{-1} - \lambda (\hat{\Sigma} + \lambda I)^{-1} \right) \right) \\
&= \frac{1}{n} \text{Tr} \left(\Sigma_t \frac{1}{t\lambda} \left(\left(\frac{1}{\lambda + t\lambda} \hat{\Sigma} + I \right)^{-1} - \left(\frac{1}{\lambda} \hat{\Sigma} + I \right)^{-1} \right) \right) \\
&= \frac{1}{n} \text{Tr} \left(\Sigma_t \frac{1}{t\lambda} \left(\left(\frac{1}{\lambda} \hat{\Sigma} + I \right)^{-1} \left(\frac{1}{\lambda} \hat{\Sigma} + I - \frac{1}{\lambda + t\lambda} \hat{\Sigma} - I \right) \left(\frac{1}{\lambda + t\lambda} \hat{\Sigma} + I \right)^{-1} \right) \right) \\
&= \frac{1}{n} \text{Tr} \left(\Sigma_t (\hat{\Sigma} + \lambda I)^{-1} \hat{\Sigma} (\hat{\Sigma} + (\lambda + t\lambda)I)^{-1} \right) \\
&= \frac{1}{n} \text{Tr} \left((\hat{\Sigma} + (\lambda + t\lambda)I)^{-1} (\hat{\Sigma} + \lambda I)^{-1} \hat{\Sigma} \Sigma_t \right),
\end{aligned}$$

where the last line follows from the cyclic property of the trace and the commutativity of $\hat{\Sigma}$, $(\hat{\Sigma} + \lambda I)^{-1}$ and $(\hat{\Sigma} + (\lambda + t\lambda)I)^{-1}$. Plugging this into the LHS of (E.46) yields

$$\begin{aligned}
&\left| V_X^1(\lambda) - \frac{1}{t\lambda} (\tilde{V}_X^1(\lambda + t\lambda) - \tilde{V}_X^1(\lambda)) \right| \\
&= \left| \frac{1}{n} \text{Tr} \left(\left((\hat{\Sigma} + \lambda I)^{-1} - (\hat{\Sigma} + (\lambda + t\lambda)I)^{-1} \right) (\hat{\Sigma} + \lambda I)^{-1} \hat{\Sigma} \Sigma_t \right) \right| \\
&= \left| \frac{t\lambda}{n} \text{Tr} \left((\hat{\Sigma} + (\lambda + t\lambda)I)^{-1} (\hat{\Sigma} + \lambda I)^{-2} \hat{\Sigma} \Sigma_t \right) \right| \\
&\leq \left\| \Sigma_t (\hat{\Sigma} + (\lambda + t\lambda)I)^{-1} (\hat{\Sigma} + \lambda I)^{-2} \right\|_2 \frac{t\lambda}{n} \text{Tr} \hat{\Sigma} = O(t\lambda^{-2}),
\end{aligned}$$

where the last line follows from the bound $\frac{1}{n} \text{Tr} \hat{\Sigma} = O(1)$, which holds due to Proposition E.2. ♣

Let us denote the zero-centered counterparts of the corresponding V_X^1 terms as

$$\begin{aligned}
V_X^0(\hat{\beta}, \beta) &:= \frac{1}{n} \text{Tr}[\hat{\Sigma}_0^+ \Sigma_t] \\
V_X^0(\lambda) &:= \frac{1}{n} \text{Tr}[(\hat{\Sigma}_0 + \lambda I)^{-2} \hat{\Sigma}_0 \Sigma_t] = \frac{d}{d\lambda} \left(\frac{\lambda}{n} \text{Tr} \left(\Sigma_t (\hat{\Sigma}_0 + \lambda I)^{-1} \right) \right), \\
\tilde{V}_X^0(\lambda) &:= \frac{\lambda}{n} \text{Tr} \left(\Sigma_t (\hat{\Sigma}_0 + \lambda I)^{-1} \right).
\end{aligned}$$

Analogously to (E.45) and (E.46), it holds that

$$\left| V_X^0(\hat{\beta}, \beta) - V_X^0(\lambda) \right| = O(\lambda), \quad \left| V_X^0(\lambda) - \frac{1}{t\lambda} (\tilde{V}_X^0(\lambda + t\lambda) - \tilde{V}_X^0(\lambda)) \right| = O(t\lambda^{-2}). \quad (\text{E.47})$$

The next step is to prove that, for $1 > \lambda > p^{-0.49}$,

$$\tilde{V}_X^1(\lambda) = \tilde{V}_X^0(\lambda) + O\left(\frac{\lambda^{-2}}{n}\right). \quad (\text{E.48})$$

Proof of the claim in (E.48). Expanding the expression, we want to prove that

$$\tilde{V}_X^1(\lambda) = \frac{\lambda}{n} \text{Tr} \left(\Sigma_t (\hat{\Sigma} + \lambda I)^{-1} \right) = \frac{\lambda}{n} \text{Tr} \left(\Sigma_t (\hat{\Sigma}_0 + \lambda I)^{-1} \right) + O\left(\frac{\lambda^{-2}}{n}\right).$$

Notice that $\tilde{V}_X^1(\lambda)$ crucially contains $(\hat{\Sigma} + \lambda I)^{-1}$ in its expression, which we have already analyzed in the context of $B_X^1(\hat{\beta}, \beta)$. Recalling the definitions of A, u, v, U, C, a, b, d , and Δ from (E.37),

(E.38), and (E.39), we can then expand $\tilde{V}_X^1(\lambda)$ as

$$\begin{aligned} \frac{\lambda}{n} \text{Tr} \left(\Sigma_t (\hat{\Sigma} + \lambda I)^{-1} \right) &= \frac{\lambda}{n} \text{Tr} \left(\Sigma_t (A + UCU^\top)^{-1} \right) \\ &= \frac{\lambda}{n} \text{Tr} \left(\Sigma_t (A^{-1} - A^{-1}U(C^{-1} - U^\top A^{-1}U)^{-1}U^\top A^{-1}) \right) \\ &= \frac{\lambda}{n} \text{Tr} (\Sigma_t A^{-1}) + \hat{S}, \end{aligned}$$

where $\hat{S} := -\frac{\lambda}{n} \text{Tr} (\Sigma_t A^{-1}U(C^{-1} - U^\top A^{-1}U)^{-1}U^\top A^{-1})$.

We now prove that $\hat{S}$ is small. To do so, we decompose

$$\begin{aligned} \hat{S} &= \frac{\lambda}{n} \text{Tr} \left(\Sigma_t \frac{1}{\Delta} A^{-1} \left(\left(n \frac{\gamma_t}{\gamma} + d \right) uu^\top + (1-b)(uv^\top + vu^\top) + a vv^\top \right) A^{-1} \right) \\ &= \hat{T}_{u,u} + \hat{T}_{u,v} + \hat{T}_{v,v}, \end{aligned}$$

where $\hat{T}_{u,u}$ is the summand corresponding to $uu^\top$, $\hat{T}_{u,v}$ to $uv^\top + vu^\top$, and $\hat{T}_{v,v}$ to $vv^\top$. Zooming in on one of the terms, it holds that

$$\begin{aligned} \hat{T}_{u,u} &= \frac{\lambda}{n} \text{Tr} \left(\Sigma_t \frac{1}{\Delta} A^{-1} \left(n \frac{\gamma_t}{\gamma} + d \right) uu^\top A^{-1} \right) \\ &= \frac{\lambda}{n} \frac{n \frac{\gamma_t}{\gamma} + d}{\Delta} \text{Tr} (\Sigma_t A^{-1} uu^\top A^{-1}) \\ &= \frac{\lambda}{n} \frac{n \frac{\gamma_t}{\gamma} + d}{\Delta} u^\top A^{-1} \Sigma_t A^{-1} u. \end{aligned}$$

Note that

$$\|A^{-1} \Sigma_t A^{-1}\|_2 \leq \frac{\lambda^{-2}}{\tau},$$

and $\|u\|_2 \leq c$. Using this, we get that, with high probability, it holds

$$|\hat{T}_{u,u}| \leq c \frac{\lambda^{-2}}{n}.$$

With analogous passages, we have that

$$|\hat{T}_{u,v}| \leq c \frac{\lambda^{-2}}{n}, \quad |\hat{T}_{v,v}| \leq c \frac{\lambda^{-2}}{n}$$

holds with high probability over the sampling of Z . Putting all together, we get

$$\frac{\lambda}{n} \text{Tr} (\Sigma_t (\hat{\Sigma} + \lambda I)^{-1}) = \frac{\lambda}{n} \text{Tr} (\Sigma_t A^{-1}) + O\left(\frac{\lambda^{-2}}{n}\right).$$

Lastly, doing all of this again to take out the terms containing μ_s from A , i.e., by taking

$$\tilde{A} = A - \frac{X^{0\top} 1_{n_s} \mu_s^\top}{n} - \frac{\mu_s 1_{n_s}^\top X^0}{n} - \frac{\gamma_s}{\gamma} \mu_s \mu_s^\top = \hat{\Sigma}_0 + \lambda I,$$

we get

$$\frac{\lambda}{n} \text{Tr} (\Sigma_t (\hat{\Sigma} + \lambda I)^{-1}) = \frac{\lambda}{n} \text{Tr} \left(\Sigma_t (\hat{\Sigma}_0 + \lambda I)^{-1} \right) + O\left(\frac{\lambda^{-2}}{n}\right),$$

proving the claim. ♣

Finally, combining (E.45), (E.46), (E.47) and (E.48), for $1 > \lambda > p^{-0.49}$ and $t > 0$, we have that

$$\begin{aligned}
\left| V_X^1(\hat{\beta}, \beta) - V_X^0(\hat{\beta}, \beta) \right| &\leq \left| V_X^1(\hat{\beta}, \beta) - V_X^1(\lambda) \right| + \left| V_X^1(\lambda) - V_X^0(\lambda) \right| + \left| V_X^0(\hat{\beta}, \beta) - V_X^0(\lambda) \right| \\
&\leq O(\lambda) + \left| \tilde{V}_X^1(\lambda) - \frac{1}{t\lambda} \left(\tilde{V}_X^1(\lambda + t\lambda) - \tilde{V}_X^1(\lambda) \right) \right| \\
&\quad + \left| \tilde{V}_X^0(\lambda) - \frac{1}{t\lambda} \left(\tilde{V}_X^0(\lambda + t\lambda) - \tilde{V}_X^0(\lambda) \right) \right| \\
&\quad + \left| \frac{1}{t\lambda} \left(\tilde{V}_X^1(\lambda + t\lambda) - \tilde{V}_X^1(\lambda) \right) - \frac{1}{t\lambda} \left(\tilde{V}_X^0(\lambda + t\lambda) - \tilde{V}_X^0(\lambda) \right) \right| \\
&\leq O(\lambda) + O\left(\frac{t}{\lambda^2}\right) + \frac{1}{t\lambda} \left| \tilde{V}_X^1(\lambda + t\lambda) - \tilde{V}_X^0(\lambda + t\lambda) \right| + \frac{1}{t\lambda} \left| \tilde{V}_X^1(\lambda) - \tilde{V}_X^0(\lambda) \right| \\
&= O(\lambda) + O(t\lambda^{-2}) + O\left(\frac{t^{-1}\lambda^{-3}}{n}\right).
\end{aligned}$$

Taking $t = \lambda^3$ and $\lambda = n^{-1/7}$, we get $\left| V_X^1(\hat{\beta}, \beta) - V_X^0(\hat{\beta}, \beta) \right| = O(n^{-1/7})$, proving the claim from (E.44). As $V_X(\hat{\beta}; \beta) = V_X^1(\hat{\beta}; \beta) + V_X^2(\hat{\beta}; \beta)$, and $V_X^2(\hat{\beta}, \beta) = O(1/p)$ by (E.43) we conclude that

$$V_X(\hat{\beta}; \beta) = \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}^+(\Sigma_t + \mu_t \mu_t^\top)] = \frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}_0^+ \Sigma_t] + O\left(p^{-1/7}\right).$$

Plugging in the expression of $\frac{\sigma^2}{n} \text{Tr}[\hat{\Sigma}_0^+ \Sigma_t]$ given in [Song et al., 2024, Theorem 4.1] yields, with high probability,

$$V_X(\hat{\beta}; \beta) = -\frac{\sigma^2}{\gamma} \int \frac{\lambda^t (a_3 \lambda^s + a_4 \lambda^t)}{(a_1 \lambda^s + a_2 \lambda^t + 1)^2} d\hat{H}_p(\lambda^s, \lambda^t) + O(p^{-c}),$$

where (a_1, a_2, a_3, a_4) is the unique solution, with a_1, a_2 positive, to (E.31). Taking the limit $p, n \rightarrow \infty$ gives the desired result for the variance term and concludes the proof.

E.8 Proof of Theorem D.2

For $\Sigma_t = I_p$ and $\Sigma_s \in \mathbb{R}_{>0}^{p \times p}$, it holds that

$$\mathcal{R}_o(\Sigma_s, I_p, \beta) = \mathcal{V}(\Sigma_s, I_p) + \mathcal{B}(\Sigma_s, I_p, \beta).$$

We analyze each of the two terms separately.

Calculating $\mathcal{B}(\Sigma_s, I_p, \beta)$. Note that $\Sigma_t = I_p$ implies $\lambda_i^t = 1$ in all the equations in (E.32). Plugging this in, one gets that the third and fourth equation in (E.32) are satisfied for $b_4 = b_2$ and $b_3 = b_1$. From the uniqueness of a solution (b_1, b_2, b_3, b_4) to the whole system of equations in (E.32), and the fact that b_3 and b_4 only show up in the mentioned third and fourth equation, we get that it must hold $b_4 = b_2$ and $b_3 = b_1$. Plugging this into the bias term we get that

$$\begin{aligned}
\mathcal{B}(\Sigma_s, I_p, \beta) &= \int \frac{b_3 \lambda^s + (b_4 + 1) \lambda^t}{(b_1 \lambda^s + b_2 \lambda^t + 1)^2} d\hat{G}_p(\lambda^s, \lambda^t) \\
&= \int \frac{b_1 \lambda^s + b_2 + 1}{(b_1 \lambda^s + b_2 + 1)^2} d\hat{G}_p(\lambda^s, \lambda^t) \\
&= \sum_{i=1}^p \frac{\langle \beta, u_i \rangle^2}{b_1 \lambda_i^s + b_2 + 1},
\end{aligned}$$

noting that $u_i \in \mathbb{R}^p$ is the eigenvector of the matrix Σ_s corresponding to the eigenvalue λ_i^s .

Recall that we have assumed in the setup of Section D that β is sampled from a sphere of constant radius, which we will denote by rS^{p-1} , i.e., $r = \|\beta\|_2$. We now prove concentration of $\mathcal{B}(\Sigma_s, I_p, \beta)$ over this sampling of β . Towards this end, we introduce a matrix $A \in \mathbb{R}^{p \times p}$ such that

$$\mathcal{B}(\Sigma_s, I_p, \beta) = \beta^\top A \beta, \quad A := \sum_{i=1}^p \frac{1}{b_1 \lambda_i^s + b_2 + 1} u_i u_i^\top.$$

Notice that first equation of (E.32) yields

$$\frac{1}{\gamma p} \sum_{i=1}^p \frac{b_1 \lambda_i^s + b_2}{b_1 \lambda_i^s + b_2 + 1} = 1,$$

which gives

$$\text{Tr}(A) = \sum_{i=1}^p \frac{1}{b_1 \lambda_i^s + b_2 + 1} = p - n.$$

Since both b_1 and b_2 are positive, as stated in Theorem D.1, it holds

$$\|A\|_2 = \lambda_1(A) = \frac{1}{b_1 \lambda_p^s + b_2 + 1} \leq 1.$$

Note that

$$\begin{aligned} \mathbb{E}_{\beta \sim rS^{p-1}} \beta^\top A \beta &= \mathbb{E} \sum_{i=1}^p \frac{\langle \beta, u_i \rangle^2}{b_1 \lambda_i^s + b_2 + 1} \\ &= \sum_{i=1}^p \frac{1}{b_1 \lambda_i^s + b_2 + 1} \mathbb{E} \langle \beta, u_i \rangle^2 \\ &= \frac{1}{p} \sum_{i=1}^p \frac{1}{b_1 \lambda_i^s + b_2 + 1} r^2 \\ &= \frac{p-n}{p} r^2. \end{aligned} \tag{E.49}$$

Furthermore, the function $\beta \rightarrow \beta^\top A \beta$ is Lipschitz over the sphere. Namely, for two vectors $\beta_1, \beta_2 \in rS^{p-1}$, it holds that

$$|\beta_1^\top A \beta_1 - \beta_2^\top A \beta_2| \leq |\beta_1^\top A(\beta_1 - \beta_2)| + |\beta_2^\top A(\beta_1 - \beta_2)| \leq 2r \|A\|_2 \|\beta_1 - \beta_2\|_2 \leq 2r \|\beta_1 - \beta_2\|_2.$$

Then, due to the concentration of Lipschitz functions over the sphere [Vershynin, 2018, Theorem 5.1.4], we get that, with overwhelming probability,

$$|\beta^\top A \beta - \mathbb{E} \beta^\top A \beta| = O(n^{-c_1}),$$

for any constant $c_1 < 1/2$. Plugging (E.49) gives

$$\mathcal{B}(\Sigma_s, I_p, \beta) = \beta^\top A \beta = \frac{p-n}{p} r^2 + O(n^{-c_1}),$$

with overwhelming probability. We can readily calculate the bias term for $\Sigma_s = I_p$:

$$\mathcal{B}(I_p, I_p, \beta) = \sum_{i=1}^p \frac{\langle \beta, u_i \rangle^2}{b_1 + b_2 + 1} = \frac{p-n}{p} r^2.$$

Thus, for any $\Sigma_s \in \mathbb{R}_{>0}^{p \times p}$, we have

$$\mathcal{B}(I_p, I_p, \beta) \leq \mathcal{B}(\Sigma_s, I_p, \beta) + O(n^{-c_1}), \tag{E.50}$$

with overwhelming probability.

Calculating $\mathcal{V}(\Sigma_s, I_p)$. Note that

$$\begin{aligned} \mathcal{V}(\Sigma_s, I_p) &= -\sigma^2 \frac{1}{\gamma} \int \frac{\lambda^t (a_3 \lambda^s + a_4 \lambda^t)}{(a_1 \lambda^s + a_2 \lambda^t + 1)^2} d\hat{H}_p(\lambda^s, \lambda^t) \\ &= -\sigma^2 \frac{1}{\gamma} \int \frac{a_3 \lambda^s + a_4}{(a_1 \lambda^s + a_2 + 1)^2} d\hat{H}_p(\lambda^s, \lambda^t) \\ &= \sigma^2 (a_1 + a_2), \end{aligned} \tag{E.51}$$

where the last equality follows from the third equation in (E.31) and the fact that $\lambda_i^t = 1$ for all $i \in [p]$. Moreover, subtracting the second from the first equation in (E.31) yields

$$0 = 1 - \frac{\gamma_s}{\gamma} - \frac{1}{\gamma p} \sum_{i=1}^p \frac{a_2}{a_1 \lambda_i^s + a_2 + 1}. \tag{E.52}$$

Analyzing just the first equation in (E.31), we get

$$\frac{1}{\gamma p} \left(p - \sum_{i=1}^p \frac{1}{a_1 \lambda_i^s + a_2 + 1} \right) = \frac{1}{\gamma p} \sum_{i=1}^p \frac{a_1 \lambda_i^s + a_2}{a_1 \lambda_i^s + a_2 + 1} = 1,$$

which gives

$$\sum_{i=1}^p \frac{1}{a_1 \lambda_i^s + a_2 + 1} = p - n.$$

Plugging this into (E.52) we get that $a_2 = \frac{\gamma t}{1-\gamma}$. Therefore, a_1 is the unique solution to

$$\sum_{i=1}^p \frac{1}{a_1 \lambda_i^s + c_2} = p - n, \quad (\text{E.53})$$

for $c_2 = \frac{\gamma t}{1-\gamma} + 1 > 0$. From (E.51), we have that $\mathcal{V}(\Sigma_s, I_p)$ only depends on Σ_s through a_1 , with which it monotonically increases. To conclude this section, we will apply the majorization argument from the proof of Theorem 3.2 with a slight modification. Almost all parts of the argument are analogous, and we restate them mainly for convenience.

Let us denote by $\vec{\lambda}^s := [\lambda_1^s, \dots, \lambda_p^s]$. Then, for fixed n, p and $\vec{\lambda}^s$, we will refer to $a_1(\vec{\lambda}^s)$ as the positive solution to (E.53). Note that from Theorem D.1 we have that this solution is unique. Consider a function $f : \mathbb{R}_{\geq 0}^p \rightarrow \mathbb{R}_{\geq 0}^p$. We call a function f *good*, if and only if

$$\sum_{i=1}^p \frac{1}{a_1(\vec{\lambda}^s) f(\vec{\lambda}^s)_i + c_2} < \sum_{i=1}^p \frac{1}{a_1(\vec{\lambda}^s) \lambda_i^s + c_2}. \quad (\text{E.54})$$

We claim that, if f is good, then

$$a_1(f(\vec{\lambda}^s)) < a_1(\vec{\lambda}^s). \quad (\text{E.55})$$

Proof of the claim. Consider a good function f . Then, we have

$$\sum_{i=1}^p \frac{1}{a_1(\vec{\lambda}^s) f(\vec{\lambda}^s)_i + c_2} < \sum_{i=1}^p \frac{1}{a_1(\vec{\lambda}^s) \lambda_i^s + c_2} = p - n.$$

Furthermore, setting $a_1 = 0$ we get

$$\begin{aligned} \sum_{i=1}^p \frac{1}{0 \cdot f(\vec{\lambda}^s)_i + c_2} &= p \frac{1}{\frac{\gamma t}{1-\gamma} + 1} \\ &= p \frac{p - n}{p - n_s} \\ &> p - n. \end{aligned}$$

By continuity, there exists $a'_1 \in (0, a_1(\vec{\lambda}^s))$ for which

$$\sum_{i=1}^p \frac{1}{a'_1 f(\vec{\lambda}^s)_i + c_2} = n - p,$$

implying $a_1(f(\vec{\lambda}^s)) = a'_1 < a_1(\vec{\lambda}^s)$, which concludes the proof. ♣

Next, for $i, j \in [p]$ s.t. $i < j$, we introduce a function $f_c^{i,j} : \mathbb{R}_{\geq 0}^p \rightarrow \mathbb{R}_{\geq 0}^p$ defined as

$$f_c^{i,j}(\vec{\lambda})_k = \begin{cases} \lambda_i^s - c & k = i, \\ \lambda_j^s + c & k = j, \\ \lambda_k^s & k \neq i, j, \end{cases}$$

where $c > 0$ is a constant. We now claim that $f_c^{i,j}$ is good for any $i, j \in [p]$ and $c > 0$, such that $\lambda_i^s > \lambda_j^s + c$.

Proof of the claim. The claim is equivalent to

$$\frac{1}{a_1(\vec{\lambda}^s)(\lambda_i^s - c) + c_2} + \frac{1}{a_1(\vec{\lambda}^s)(\lambda_j^s + c) + c_2} < \frac{1}{a_1(\vec{\lambda}^s)\lambda_i^s + c_2} + \frac{1}{a_1(\vec{\lambda}^s)\lambda_j^s + c_2}.$$

For simplicity, let us denote $a := a_1(\vec{\lambda}^s)$. Then,

$$\begin{aligned}
& \frac{1}{a(\lambda_i^s - c) + c_2} + \frac{1}{a(\lambda_j^s + c) + c_2} < \frac{1}{a\lambda_i^s + c_2} + \frac{1}{a\lambda_j^s + c_2} \\
& \iff \frac{a(\lambda_i^s + \lambda_j^s) + 2c_2}{(\lambda_i^s a - ca + c_2)(\lambda_j^s a + ca + c_2)} < \frac{a(\lambda_i^s + \lambda_j^s) + 2c_2}{(\lambda_i^s a + c_2)(\lambda_j^s a + c_2)} \\
& \iff (\lambda_i^s a + c_2)(\lambda_j^s a + c_2) < (\lambda_i^s a - ca + c_2)(\lambda_j^s a + ca + c_2) \\
& \iff ca(\lambda_i^s a + c_2) - ca(\lambda_j^s a + c_2) - c^2 a^2 > 0 \\
& \iff ca^2(\lambda_i^s - \lambda_j^s) > c^2 a^2 \\
& \iff \lambda_i^s > \lambda_j^s + c,
\end{aligned}$$

which proves the claim. ♣

This implies that, for $t \in (0, 1)$, transformations of the form

$$(\lambda_i^s, \lambda_j^s) \rightarrow (t\lambda_i^s + (1-t)\lambda_j^s, (1-t)\lambda_i^s + t\lambda_j^s) \quad (\text{E.56})$$

are good. Let us denote by $\vec{\lambda}^{id} := [1, \dots, 1]$, which corresponds to the matrix I_p . Pick any $\vec{\lambda}^s \neq \vec{\lambda}^{id}$ that corresponds to some matrix $\Sigma_s \in \mathcal{S}$, so it satisfies $\lambda_1^s \geq \lambda_2^s \geq \dots \geq \lambda_p^s$, as well as $\sum_{i=1}^p \lambda_i^s = p$.

Firstly, we claim that $\vec{\lambda}^{id}$ is majorized by $\vec{\lambda}^s$. Suppose otherwise, that for some $k \in [p]$

$$\sum_{i=1}^k \lambda_i^s < \sum_{i=1}^k 1 = k,$$

implying also that $\lambda_k^s < 1$. Then, we have

$$p = \sum_{i=1}^p \lambda_i^s < (p-k)\lambda_k^s + k < (p-k)1 + k = p,$$

which is a contradiction.

Next, as $\vec{\lambda}^{id}$ is majorized by $\vec{\lambda}^s$, $\vec{\lambda}^{id}$ can be derived from $\vec{\lambda}^s$ by a finite sequence of steps of the form in (E.56) with $t \in [0, 1]$, see [Marshall et al., 1979, Chapter 4, Proposition A.1]. Since both vectors $\vec{\lambda}^{id}$ and $\vec{\lambda}^s$ are non-increasing, the $t = 0$ transformation can always be omitted. Moreover, $t = 1$ is just the identity transformation, so it can also be omitted and we actually have $t \in (0, 1)$. In formulas, we have that

$$\vec{\lambda}^{id} = f_{c_l}^{i_l, j_l}(\dots f_{c_1}^{i_1, j_1}(\vec{\lambda}^s) \dots).$$

Since each of the functions above is good, we have that $a_1(\vec{\lambda}^{id}) < a_1(\vec{\lambda}^s)$. As $\mathcal{V}(\Sigma_s, I_p)$ is increasing with a_1 , this directly implies that, for any $\Sigma_s \in \mathbb{R}_{>0}^{p \times p}$,

$$\mathcal{V}(I_p, I_p) \leq \mathcal{V}(\Sigma_s, I_p).$$

Combining this with (E.50), we get

$$\mathcal{R}_o(I_p, I_p, \beta) \leq \mathcal{R}_o(\Sigma_s, I_p, \beta) + o(1),$$

with overwhelming probability, which concludes the proof.

F Additional numerical results

Setup details. We train for 200 epochs using SGD as optimizer, and we use cosine annealing; the initial learning rate is 0.1 for *Scratch* (0.2 for the experiment of Table 2a) and 0.01 for *Distillation* and *Pretrained*. The *Distillation* teacher is a ResNet-50 trained on CIFAR-10. We use an early stopping with patience 20 based on a validation subset (10% of the full training dataset). We avoid up-scaling images in the *Pretrained* experiments to better demonstrate the effect of synthetic data augmentation. On the generation side, to generate the images by T2I models, we use CLIP’s text encoder prompt template on CIFAR-10 and ImageNet labels. Moreover, as models like StableDiffusion 1.4 sometimes generate low quality data or images discarded by the safety checker, before applying all the algorithms, we do an initial pruning of 2% of the generated pool based on the distance to the CLIP embedding of the label. For RxRx1, we used images from four common perturbations (1108, 1124, 1137, 1138) on HUVEC cells as the real training samples and train a linear classifier on frozen features from an ImageNet-pretrained ResNet. For each class, MorphGen [Demirel et al., 2025] generates a pool of 500 synthetic images; we augment the real training set (30 images/class) with 60 selected synthetic images/class and evaluate on a disjoint test set of 20 images/class. We repeat the experiment 10 times by resampling the real subset from 120 images/class. As in the main setup, CLIP features are used for the selection algorithms.

Table 2: *Covariance matching* performs on par with the best baselines for two additional datasets. In (a), we train a ResNet-18 from scratch on ImageNet-100 with synthetic images from StyleGAN-XL and T2I models. In (b), we train a linear model on top of an ImageNet-pretrained ResNet for perturbation classification on a small subset of RxRx1 [Sypetkowski et al., 2023] augmented with synthetic images from MorphGen [Demirel et al., 2025].

Method	Truncated models	T2I models	Method	MorphGen
No synthetic	40.78 $\pm$ 1.29		No synthetic	86.83 $\pm$ 2.44
Center matching [He et al., 2023]	53.39 $\pm$ 0.37	53.96 $\pm$ 1.06	Center matching [He et al., 2023]	88.17 $\pm$ 2.35
DS3 [Hulkund et al., 2025]	57.47 $\pm$ 0.87	53.51 $\pm$ 0.31	Random	87.33 $\pm$ 2.03
Random	54.14 $\pm$ 0.82	49.84 $\pm$ 1.32	K-means [Lin et al., 2023]	89.00 $\pm$ 1.70
Text matching [Lin et al., 2023]	53.39 $\pm$ 0.99	53.37 $\pm$ 0.72	DS3 [Hulkund et al., 2025]	89.67 $\pm$ 1.45
Covariance matching (ours)	57.52 $\pm$ 0.36	53.07 $\pm$ 0.89	Center sampling [Lin et al., 2023]	88.75 $\pm$ 2.27
Real upper bound	62.67 $\pm$ 0.65		Covariance matching (ours)	90.00 $\pm$ 1.86

(a) ImageNet-100 dataset

(b) RxRx1 dataset

Baseline Details. We briefly describe all baselines compared against *Covariance matching* in the main text.

- *Center matching* [He et al., 2023]: selects the n_s generated images nearest to the centroid of the n_t real training features.
- *Center sampling* [Lin et al., 2023]: samples generated images with probability proportional to their cosine similarity to the n_t real training features.
- *DS3* [Hulkund et al., 2025]: clusters the generated pool into 200 clusters; for each real image, retains its nearest cluster and uniformly samples n_s images from the retained set.
- *K-means* [Lin et al., 2023]: clusters the generated pool into n_s clusters and selects one representative per cluster.
- *Random*: uniformly samples n_s images from the generated pool. The methods “No-filtering,” “Match-dist,” and “Match-label” [Hulkund et al., 2025] are equivalent to this baseline in our setting, since each class has the same number of data points.
- *Text matching* [Lin et al., 2023]: selects the n_s generated images nearest to the class text embedding.
- *Text sampling* [Lin et al., 2023]: samples generated images with probability proportional to their cosine similarity to the class text embedding.
- *No synthetic*: uses only n_t real samples from the training distribution (synthetic data discarded).
- *Real upper bound*: uses $n_t + n_s$ real samples from the training distribution (synthetic data replaced by in-distribution data).

F.1 Controlled experiments

In Table 3, we consider *zero-diversity generators*. Specifically, for each class, we combine 2K StyleGAN2-Ada images with a total of 8K images produced by two zero-diversity generators. Each of these generators emits a single prototype per class: one near the class center of the real samples, and one near the class label’s CLIP embedding. This yields high precision, but low diversity relative to the real distribution. Our results show that, again, covariance matching performs well as it avoids selecting many samples with low diversity (collapsed clusters). In contrast, not fully taking into account the diversity of selected samples, methods like DS3 perform rather poorly. In Figure 2, we consider *inserting images from the target distribution* into the pool of synthetic images and test the ability of different methods to select them. Specifically, we form a pool of 4K StableDiffusion1.4 images and 1K images from the target distribution (different from the $n_t = 200$ images forming the training distribution), letting each method take $n_s = 800$. Our results show that covariance matching selects the highest fraction of images coming from the target distribution, whereas other selectors largely fail to do so.

Zero-diversity generators. To assess the importance of filtering low-diversity data, we construct a pool per CIFAR-10 class with 2K images from StyleGAN2-Ada and 8K images from two collapsed generators. The first collapsed model emits the image whose CLIP embedding is closest to the class label; the second produces images near the mean embedding of the class’s real subset. We sample 4K images from each collapsed generator, yielding a total 10K images per class. As shown in Table 3, most baselines over-select from the collapsed generators because they ignore the diversity of selected samples. In particular, DS3 retains the two clusters formed by the collapsed outputs and thus fails to filter them. By contrast, K-means and Covariance matching draw more from the 2K non-collapsed subset and achieve higher classification accuracy.

Table 3: *Covariance matching* performs on par with the best baselines across three training paradigms on CIFAR-10, when the synthetic data is generated via a StyleGAN2-Ada model and two zero-diversity generators.

Method	Scratch	Distillation	Pretrained
No synthetic	44.36 $\pm$ 1.51	47.33 $\pm$ 0.57	63.40 $\pm$ 1.33
Center matching [He et al., 2023]	45.33 $\pm$ 2.43	47.50 $\pm$ 0.55	62.96 $\pm$ 1.26
Center sampling [Lin et al., 2023]	46.88 $\pm$ 2.59	51.11 $\pm$ 0.60	65.38 $\pm$ 1.14
DS3 [Hulkund et al., 2025]	53.74 $\pm$ 1.92	59.16 $\pm$ 1.56	69.43 $\pm$ 0.93
K-means [Lin et al., 2023]	60.20 $\pm$ 1.35	65.03 $\pm$ 0.81	72.83 $\pm$ 0.48
Random	50.31 $\pm$ 1.28	51.82 $\pm$ 0.91	66.27 $\pm$ 1.21
Text matching [Lin et al., 2023]	42.89 $\pm$ 1.89	47.38 $\pm$ 0.76	62.82 $\pm$ 1.31
Text sampling [Lin et al., 2023]	48.13 $\pm$ 1.81	50.81 $\pm$ 0.77	66.12 $\pm$ 1.06
Covariance matching (ours)	58.97 $\pm$ 1.67	64.85 $\pm$ 0.63	72.38 $\pm$ 0.66
Real upper bound	61.08 $\pm$ 2.54	65.38 $\pm$ 0.51	74.35 $\pm$ 0.56

Leak experiment. We consider inserting (“leaking”) images from the target distribution into the pool of synthetic images and test the ability of different methods to select them. We use 1K leaked CIFAR-10 images, disjoint from the 200 (n_t) real reference samples. From a pool of 4K StableDiffusion1.4 images and 1K leaked images, each method selects 800 (n_s). Figure 2 shows, for each method, the fraction of selected samples drawn from the leak. Because replacing synthetic with real augmentations yields the best accuracy (*Real upper bound*), an effective selector should prioritize leaked real images: covariance matching does, achieving the highest leaked fraction among all methods.

F.2 Ablations

In Tables 5-6, we repeat the experiments of Table 1 with DINO instead of CLIP features, demonstrating that the gains of covariance matching are not tied to a particular feature extractor. In Table 7, we compare covariance matching with the direct optimization of the objective given by Theorem 3.1. As the outcomes of these two procedures are largely similar, this further justifies the covariance matching objective. In Table 8, we show that our findings replicate in an over-parameterized regime. Finally, in Table 9, we examine the distribution of selections produced by each method, quantifying alignment with the test distribution and identifying which metrics best predict downstream accuracy.

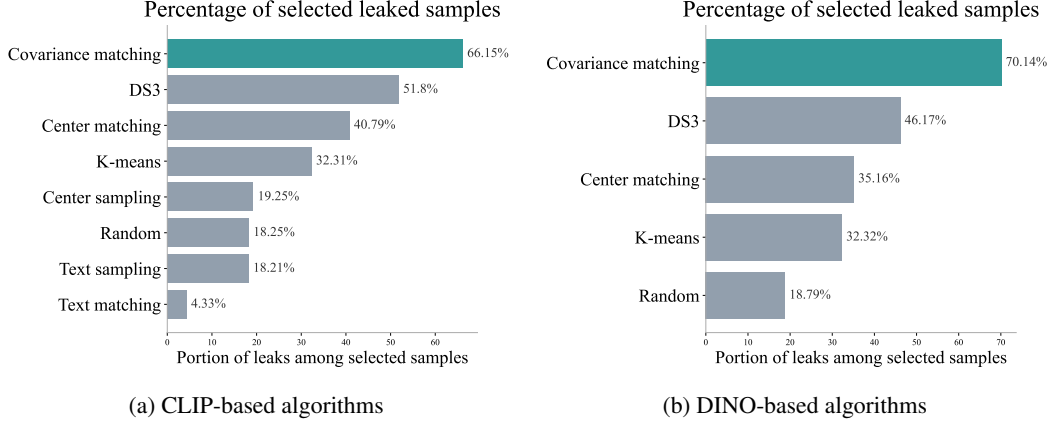


Figure 2: The portion of samples chosen from the set of leaked images shows that our proposed algorithm reliably selects real samples among the pool of generated examples.

Transformer-based models. In Table 4, we use the same setup as Table 1, but instead of ResNet, we train a ViT and a Swin-T model from scratch. We use a patch size of 4 and Adam optimizer with learning rate 0.0001 for this experiment. We observe that, in accordance with our previous findings, covariance matching surpasses other algorithms.

Table 4: *Covariance matching* outperforms all baselines when fully training a transformer model on a mix of real and synthetic data.

Method	ViT		Swin-T	
	Scratch	Distillation	Scratch	Distillation
No synthetic	40.11 $\pm$ 0.59	40.32 $\pm$ 1.01	40.02 $\pm$ 0.70	40.84 $\pm$ 0.73
Center matching [He et al., 2023]	43.89 $\pm$ 0.97	45.61 $\pm$ 0.68	44.39 $\pm$ 0.54	46.64 $\pm$ 0.53
Center sampling [Lin et al., 2023]	43.89 $\pm$ 0.95	46.29 $\pm$ 0.80	43.94 $\pm$ 1.76	46.97 $\pm$ 0.59
DS3 [Hulkund et al., 2025]	45.92 $\pm$ 0.49	48.61 $\pm$ 0.67	46.57 $\pm$ 0.68	49.55 $\pm$ 0.72
K-means [Lin et al., 2023]	44.24 $\pm$ 1.13	47.44 $\pm$ 0.97	44.71 $\pm$ 0.32	48.49 $\pm$ 0.64
Random	44.07 $\pm$ 0.82	46.50 $\pm$ 0.78	44.38 $\pm$ 0.77	47.35 $\pm$ 0.50
Text matching [Lin et al., 2023]	44.57 $\pm$ 0.57	46.02 $\pm$ 1.00	45.15 $\pm$ 0.58	46.55 $\pm$ 2.52
Text sampling [Lin et al., 2023]	43.80 $\pm$ 0.98	46.00 $\pm$ 0.98	44.59 $\pm$ 0.93	47.62 $\pm$ 0.71
Covariance matching (ours)	46.09 $\pm$ 0.91	49.53 $\pm$ 0.61	46.64 $\pm$ 0.96	50.73 $\pm$ 0.44
Real upper bound	51.85 $\pm$ 0.47	53.11 $\pm$ 0.43	52.43 $\pm$ 1.39	54.80 $\pm$ 0.69

Changing the feature extractor. In the main experiments, we use CLIP features for all selection methods. To test the dependence on the feature extractor, we repeat the setup of Table 1 with DINO-v2 features. As shown in Tables 5-6, covariance matching matches or surpasses the best baseline across settings, indicating that its effectiveness is not tied to a specific feature extractor. We also repeat the leak experiment of Figure 2, see the bar plot in (b), showing again similar results.

Table 5: *Covariance matching* outperforms all baselines across three training paradigms on CIFAR-10, when the synthetic data is generated via truncated generative models and features are extracted with DINO-v2.

Method	Scratch	Distillation	Pretrained
No synthetic	44.36 $\pm$ 1.51	47.33 $\pm$ 0.57	63.40 $\pm$ 1.33
Center matching [He et al., 2023]	50.06 $\pm$ 1.45	54.50 $\pm$ 0.62	66.23 $\pm$ 0.72
DS3 [Hulkund et al., 2025]	52.93 $\pm$ 1.65	58.69 $\pm$ 0.81	68.04 $\pm$ 0.71
K-means [Lin et al., 2023]	51.66 $\pm$ 2.10	55.97 $\pm$ 0.58	67.00 $\pm$ 0.84
Random	49.97 $\pm$ 2.45	54.79 $\pm$ 0.68	66.57 $\pm$ 0.92
Text matching [Lin et al., 2023]	51.52 $\pm$ 1.67	55.17 $\pm$ 0.57	67.13 $\pm$ 0.45
Covariance matching (ours)	54.97 $\pm$ 2.60	59.41 $\pm$ 0.81	68.87 $\pm$ 0.41
Real upper bound	61.08 $\pm$ 2.54	65.38 $\pm$ 0.51	74.35 $\pm$ 0.56

Table 6: *Covariance matching* performs on par with the best baseline across three training paradigms on CIFAR-10, when the synthetic data is generated via text-to-image (T2I) generative models and features are extracted with DINO-v2.

Method	Scratch	Distillation	Pretrained
No synthetic	44.36 ± 1.51	47.33 ± 0.57	63.40 ± 1.33
Center matching [He et al., 2023]	51.75 ± 2.01	55.67 ± 0.63	66.00 ± 0.58
DS3 [Hulkund et al., 2025]	52.33 ± 2.07	58.80 ± 0.96	66.68 ± 0.63
K-means [Lin et al., 2023]	51.14 ± 1.90	56.93 ± 0.46	65.71 ± 0.71
Random	50.45 ± 1.41	55.86 ± 0.73	65.67 ± 0.82
Text matching [Lin et al., 2023]	51.38 ± 1.51	55.81 ± 0.65	65.76 ± 1.00
Covariance matching (ours)	52.65 ± 1.47	58.78 ± 0.53	67.04 ± 0.83
Real upper bound	61.08 ± 2.54	65.38 ± 0.51	74.35 ± 0.56

Optimizing the theoretical objective. We also implement a greedy algorithm that, at each step, adds the sample minimizing the objective in (3.1) (*Alpha matching*). This method requires computing the eigenvalues of the current sample covariance and is therefore more costly than *Covariance matching*. As in *Covariance matching*, we first fit PCA on the real samples and project all features, then iteratively add the sample that yields the smallest value of (3.1). Without loss of generality, we drop the noise variance term since it scales all candidates equally. The results of Table 7 show that *Alpha matching* performs similarly to *Covariance matching*.

Table 7: *Covariance matching* performs on par with *Alpha matching* across the experiments on CIFAR-10.

Experiment	Method	Scratch	Distillation	Pretrained
Zero-diversity models	Covariance matching	58.97 ± 1.67	64.85 ± 0.63	72.38 ± 0.66
	Alpha matching	59.30 ± 2.50	64.72 ± 0.55	72.76 ± 0.73
Truncated models	Covariance matching	54.00 ± 1.89	59.77 ± 0.61	69.20 ± 0.56
	Alpha matching	52.25 ± 2.11	59.18 ± 0.68	68.32 ± 0.58
T2I models	Covariance matching	54.45 ± 2.11	59.17 ± 0.64	66.69 ± 0.70
	Alpha matching	53.37 ± 1.85	59.03 ± 0.64	66.23 ± 0.66

Over-parameterized setting. We repeat the setup of Table 1 taking $n_s = 200$ (instead of $n_s = 800$). This gives a total of $n_s + n_t = 400$ samples, which is less than the number of features $p = 512$, thus placing us in an over-parameterized regime. As shown in Table 8, the quantitative trends mirror those in the under-parameterized case.

Table 8: *Covariance matching* outperforms all baselines across three training paradigms on CIFAR-10, when the synthetic data is generated via truncated StyleGAN2-Ada models [Karras et al., 2019] in the over-parameterized regime with 200 training and 200 augmenting synthetic samples.

Method	Scratch	Distillation	Pretrained
No synthetic	44.36 ± 1.51	47.33 ± 0.57	63.40 ± 1.33
Center matching [He et al., 2023]	46.45 ± 1.97	50.83 ± 0.50	64.40 ± 1.11
Center sampling [Lin et al., 2023]	47.29 ± 1.33	50.89 ± 0.78	65.64 ± 0.74
DS3 [Hulkund et al., 2025]	48.09 ± 2.04	52.65 ± 0.61	66.41 ± 1.35
K-means [Lin et al., 2023]	47.75 ± 0.82	51.56 ± 0.68	65.47 ± 0.99
Random	47.39 ± 1.63	50.96 ± 0.22	65.49 ± 1.12
Text matching [Lin et al., 2023]	47.56 ± 1.09	51.67 ± 0.65	65.74 ± 0.78
Text sampling [Lin et al., 2023]	46.93 ± 1.95	50.64 ± 0.49	65.13 ± 1.13
Covariance matching (ours)	48.95 ± 1.28	53.28 ± 0.45	66.62 ± 0.57
Real upper bound	50.79 ± 1.70	54.66 ± 0.91	68.97 ± 0.88

Distribution of selected samples. Beyond accuracy, we assess how well each method’s selections match the test distribution. In the CIFAR-10 setup of Table 1, each method selects 800 samples per class given 200 real samples. We then calculate how well these samples match the CIFAR-10 training

dataset. The selection obtained via *Covariance matching* consistently achieves lower FID/KID and covariance distance than all other baselines. Metrics that couple fidelity and diversity (e.g., FID/KID) show larger gains than quality metrics (e.g., Precision [Kynkäänniemi et al., 2019], Density [Naeem et al., 2020]), indicating improved distributional alignment rather than mere sample quality. The results are reported in Table 9.

Table 9: *Covariance matching* selects samples that better match the target distribution according to various evaluation metrics.

Method	FID ↓	KID ↓	Precision ↑	Recall ↑	Density ↑	Coverage ↑	Covariance Shift ↓
K-means [Lin et al., 2023]	366.52 ± 2.62	0.59 ± 0.04	0.77 ± 0.01	0.41 ± 0.00	0.87 ± 0.04	0.58 ± 0.01	118.91 ± 0.62
Center matching [He et al., 2023]	544.56 ± 5.57	0.83 ± 0.06	0.78 ± 0.01	0.33 ± 0.01	0.82 ± 0.03	0.49 ± 0.01	212.55 ± 3.03
Center sampling [Lin et al., 2023]	450.27 ± 3.86	0.61 ± 0.04	0.77 ± 0.01	0.44 ± 0.01	0.86 ± 0.03	0.53 ± 0.01	150.49 ± 0.79
DS3 [Hulkund et al., 2025]	273.59 ± 6.72	0.42 ± 0.04	0.79 ± 0.01	0.45 ± 0.01	0.84 ± 0.03	0.64 ± 0.01	106.52 ± 2.44
Random	458.39 ± 4.16	0.63 ± 0.04	0.77 ± 0.02	0.44 ± 0.01	0.86 ± 0.05	0.53 ± 0.01	150.66 ± 1.08
Text matching [Lin et al., 2023]	454.23 ± 2.66	0.69 ± 0.05	0.81 ± 0.01	0.36 ± 0.00	0.90 ± 0.03	0.54 ± 0.01	172.70 ± 0.66
Text sampling [Lin et al., 2023]	447.53 ± 3.99	0.61 ± 0.04	0.77 ± 0.01	0.44 ± 0.01	0.86 ± 0.03	0.53 ± 0.01	149.98 ± 0.95
Covariance matching (ours)	242.09 ± 1.93	0.41 ± 0.04	0.78 ± 0.01	0.50 ± 0.01	0.84 ± 0.03	0.68 ± 0.01	95.55 ± 0.58

G Future work

Future work could extend the analysis to multiple Gaussian mixtures, which corresponds to optimizing the actual risk as opposed to modeling individual classes. We speculate that this may yield different insights when the training data have extremely imbalanced or fine-grained classes. It would also be interesting to introduce a model shift (different β between synthetic and real samples). In fact, synthetic data often has small differences compared to real data, which a model may overfit on, and the phenomenon could be the cause of the collapse sometimes observed in practice [Shumailov et al., 2024]. Finally, we have only focused on generalization, but other quantities may be studied in this framework, including uncertainty calibration [Nixon et al., 2019], differential privacy [Dwork, 2006], fairness [Barocas et al., 2020], and validity for prediction-powered causal inference [Cadei et al., 2025].