

THE CURSE OF OVERPARAMETRIZATION IN ADVERSARIAL TRAINING: PRECISE ANALYSIS OF ROBUST GENERALIZATION FOR RANDOM FEATURES REGRESSION

BY HAMED HASSANI^{1,a}, ADEL JAVANMARD^{2,b} 

¹*Department of Electrical and Systems Engineering, University of Pennsylvania, hassani@seas.upenn.edu*

²*Data Sciences and Operations Department, University of Southern California, ajavanma@usc.edu*

Successful deep learning models often involve training neural network architectures that contain more parameters than the number of training samples. Such overparametrized models have recently been extensively studied, and the virtues of overparametrization have been established from both the statistical perspective, via the double-descent phenomenon, and the computational perspective via the structural properties of the optimization landscape. Despite this success, it is also well known that these models are highly vulnerable to small adversarial perturbations in their inputs. Even when adversarially trained, their performance on perturbed inputs (robust generalization) is considerably worse than their best attainable performance on benign inputs (standard generalization). It is thus imperative to understand how overparametrization fundamentally affects robustness.

In this paper, we will provide a precise characterization of the role of overparametrization on robustness by focusing on random features regression models (two-layer neural networks with random first layer weights). We consider a regime where the sample size, the input dimension and the number of parameters grow proportionally, and derive an asymptotically exact formula for the robust generalization error when the model is adversarially trained. Our developed theory reveals the nontrivial effect of overparametrization on robustness and indicates that high overparametrization can hurt robust generalization.

1. Introduction The success of deep learning models is often reliant on training highly complex neural networks whose number of parameters is much larger than the number of data points. Even though the large complexity of such models allows for perfect interpolation of the data, they often achieve low generalization error. This behavior has resulted in a growing body of work aiming to analyze such so-called overparametrized models.

Recent work has demonstrated the virtues of overparametrization from statistical and optimization-based perspectives. From the statistical viewpoint, it is now well-documented that many overparametrized models exhibit a ‘double-descent’ property [5, 4, 61]: As the model complexity increases, the generalization error first follows the traditional U-shaped curve until a specific point, after which the error decreases, and attains a global minimum in the overparametrized regime. In fact, the minimum generalization error often appears to be at infinite complexity – the more overparametrized is the model, the smaller is the error. It is often argued that the good generalization behavior of highly overparametrized models is due to the inductive bias of gradient-based algorithms which helps with selecting models that generalize well –see e.g. [3, 37, 79, 35]. From the optimization viewpoint, training deep neural networks in general involves optimizing highly non-convex functions, but it has been conjectured that

MSC2020 subject classifications: Primary 62E20, 62F12; secondary 62F35.

Keywords and phrases: adversarial training, random features models, precise high-dimensional asymptotics, Gaussian equivalence property.

highly overparametrized models are easy to optimize despite non-convexity. Instances of this observation has been formally proved, e.g. in [77, 44, 41, 3, 64]. The high-level intuition here is that in the highly overparametrized regimes, a model that perfectly interpolates the training data (and so is a global minimizer of the empirical risk) is found in the neighborhood of most initializations.

Despite the remarkable success of deep neural networks, and the crucial role of overparametrization in both the generalization and the tractability aspects, these models are known to be highly vulnerable to perturbations in the input [6, 84]. With an unguarded training approach, these models show unsatisfactory *robust generalization error* in the presence of “small” worst-case perturbations to their inputs, a.k.a *adversarial examples*. This suggests that learning algorithms, even those with excellent performance on test data, may not be learning the true underlying concepts that determine the response; although they work well on naturally occurring data, adversarial examples have low probability in the data distribution and expose fundamental blind spots in the learning algorithms.

This observation stimulated significant effort to improve robustness using a wide variety of *adversarial training* methods which often involve augmenting the training loss so as to become more robust to input perturbations (see e.g. [32, 47, 43, 57, 93, 97, 13]). However, there is still a large gap between the robust generalization error and the (standard) generalization error in adversarially-trained models. In summary, while modern overparametrized machine learning models perform very well on benign inputs, they still remain fragile to perturbations in the input. These findings raise a fundamental question:

How does overparametrization affect robustness to perturbations in the input?

A few recent papers have begun to answer the above question in specific settings with rather conflicting messages:

- [46] and [23] have studied high-dimensional *linear* models and showed that the robust generalization error of adversarially-trained models becomes *worse* as the models become more overparametrized. It should be noted that for linear models, even in the case where there is no adversary, the best generalization error is attained when the model is underparametrized [37].
- Another line of work provably shows that in order to *interpolate* the training data smoothly, while being robust, overparametrization is *necessary* [10, 9]. However, we note that, in order to train robust models, it may not be beneficial to interpolate the training data as robustness is measured with worst-case performance over all the points in a neighborhood around the input data. Indeed, [22] study the tradeoffs between memorization (of training data) and robustness of two-layer neural networks and established a lower-bound on the non-robustness of the model (via the Sobolev-seminorm of the model) as an increasing function of the amount of memorization.

We will provide a more detailed discussion of these points and other related works in Section 3. Despite such interesting recent progress, a comprehensive understanding on how overparametrization precisely affects robustness remains largely mysterious.

In this paper, we focus on random features regression models that are adversarially trained using robust empirical risk minimization and provide a “precise characterization” of the robust generalization. Our analysis is carried out in a high-dimensional regime where the size of the training data n , the number of parameters N , and the dimension of the data d grow proportional to each other, i.e. $N/d \rightarrow \psi_1$ and $n/d \rightarrow \psi_2$. We further assume that the perturbations are bounded in terms of ℓ_2 norm by a value $\varepsilon > 0$. Our developed theory allows us to precisely characterize the effect of overparamterization on model robustness. One of the

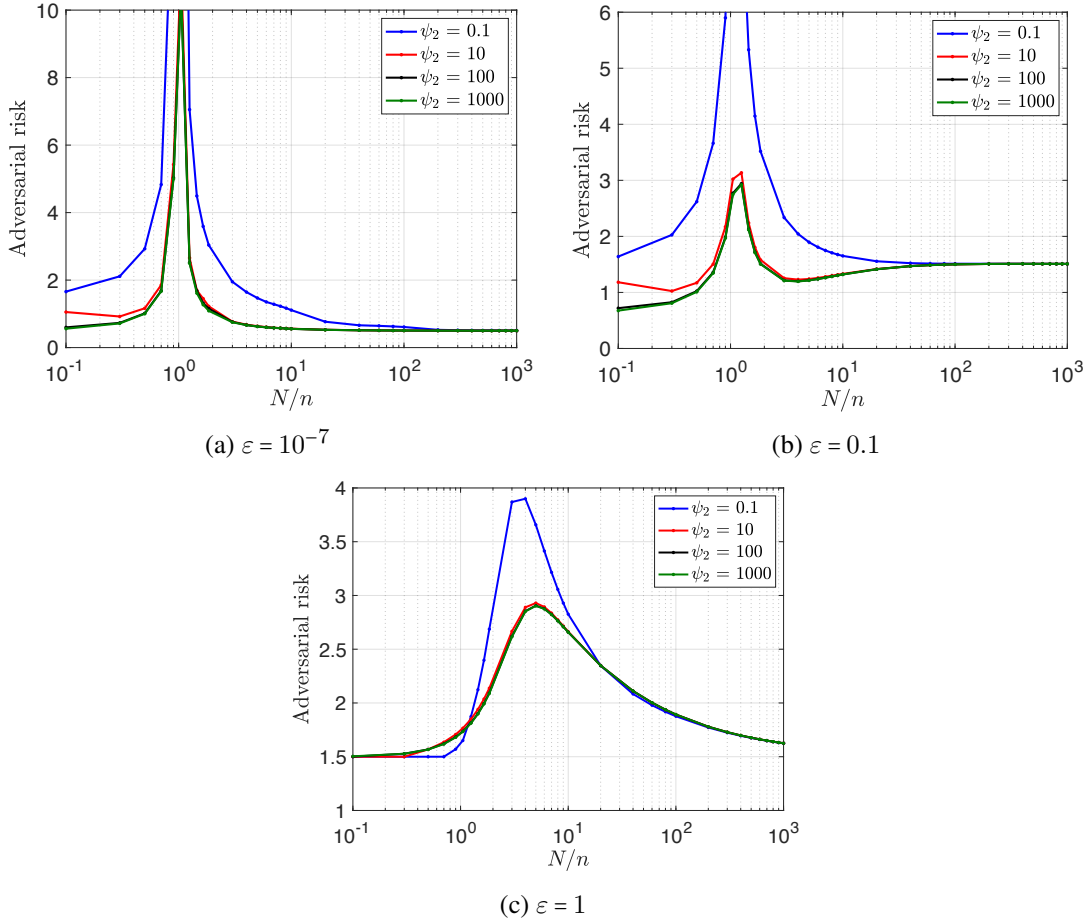


Fig 1: Random features regression with (shifted) ReLU activation ($\sigma(x) = \max(x, 0) - 1/\sqrt{2\pi}$). Data $(\mathbf{x}_i, y_i)$ is generated with d -dimensional normal covariates $\mathbf{x}_i$ and $y_i = \beta^\top \mathbf{x}_i + \xi_i$, where the noise variables $\xi_i \sim \mathcal{N}(0, 0.5)$ and $\|\beta\|_{\ell_2} = 1$. Perturbations are allowed within an Euclidean ball of radius ε , and the models are adversarially trained. We plot the robust generalization error (using Theorem 4.2) versus the amount of overparametrization N/n , where N is the number of parameters and n is the number of training data points. The plots are obtained for different values of ε and $\psi_2 = n/d$.

main consequences of our analysis is that, in general, higher overparametrization leads to a *worse* robust generalization error for the adversarially-trained models. Figure 1 depicts how the robust generalization error varies with respect to the amount of overparametrization N/n . The left figure corresponds to the case where there is no adversary (i.e. $\varepsilon \approx 0$). In this case, the robust generalization coincides with the (standard) generalization error, and is minimized at infinite overparametrization. However, as seen in the other two figures (for $\varepsilon > 0$), overparametrization is in general hurting robustness. This is clearly seen in Figure 1(c) and (b) (for $\psi_2 \geq 10$) where the minimum robust error is attained when the model is underparametrized. We refer to Figure 4 for an extended version of Figure 1 with more choices of ε and signal-to-noise ratios.

We proceed by providing an informal overview of our results and their implications in Section 2. Related works are discussed in Section 3. The main result of the paper, which characterizes the robust generalization error for the random features model is explained in Section 4. The architecture of the proof of the main result is sketched in Section 6. Our analysis develops a set of techniques that are of independent interest: (i) We derive an asymptotic

closed form for adversarial examples in trained random features models; (ii) While features are highly non-Gaussian in random features models, we prove a Gaussian equivalence property which relates robust generalization in these models to that of linear models with Gaussian features under the same correlation structure; (iii) Our analysis of the equivalent Gaussian model relies on the Convex Gaussian Min-max Theorem, which is a generalized and tight version of Gordon’s Gaussian comparison inequalities.

2. Results and discussion: An informal overview

Problem setting. Consider a supervised learning scenario where we are given i.i.d data $\{(\mathbf{x}_i, y_i)\}_{i \leq n}$ generated according to the following distribution:

$$(2.1) \quad y_i = \langle \mathbf{x}_i, \boldsymbol{\beta} \rangle + \xi_i, \quad \text{with} \quad \mathbf{x}_i \sim_{iid} \mathcal{N}(0, \mathbf{I}_d), \quad \xi_i \sim \mathcal{N}(0, \tau^2).$$

The (linear) dependence between $(\mathbf{x}_i, y_i)$ is unknown and the goal is to fit a model to this data which can be then used to predict labels for the unlabeled examples at test time.

We consider modeling the relation between label y and feature vector $\mathbf{x}$ using the class of random features (RF) model, which can be described as

$$(2.2) \quad \mathcal{F}_{\text{RF}}(\mathbf{W}) = \left\{ f(\mathbf{x}, \boldsymbol{\theta}, \mathbf{W}) = \sum_{\ell=1}^N \theta_{\ell} \sigma(\langle \mathbf{w}_{\ell}, \mathbf{x} \rangle) : \boldsymbol{\theta} = (\theta_1, \dots, \theta_N) \in \mathbb{R}^N \right\},$$

where $\boldsymbol{\theta}$ is the parameter vector to be learned and $\mathbf{W} \in \mathbb{R}^{N \times d}$ is a fixed matrix whose rows $\mathbf{w}_{\ell}$ are chosen randomly and independently of data. For simplicity we assume the normalization $\|\mathbf{w}_{\ell}\|_{\ell_2} = 1$. Namely, the vectors $\mathbf{w}_{\ell}$ are chosen uniformly at random from the unit sphere, $\mathbf{w}_{\ell} \sim \text{Unif}(\mathbb{S}^{d-1})$, which implies that $\langle \mathbf{w}_{\ell}, \mathbf{x}_j \rangle$ is of order one. In addition, $\sigma : \mathbb{R} \mapsto \mathbb{R}$ is a nonlinear activation function.

Note that in random features model training is only done on $\boldsymbol{\theta}$ and not on $\mathbf{W}$. In other words, the random features model can be perceived as a two-layer neural network with the weights of the first layer chosen randomly and independently from data, while the weights in the second layer are learned during the training phase. The random features model was introduced by [68] for scaling kernel methods to large datasets, and there has been a large body of work drawing connections between random features models, kernel methods and fully trained neural networks [15, 14, 42, 51]. The random features models are arguably the simplest analytically tractable models that capture all the features of the double descent phenomenon without assuming ad hoc misspecification structures [63]. In particular, they allow to disentangle the number of parameters from the covariates dimension and hence isolate the effects of overparametrization from the effects of the ambient dimension.

To quantify robust generalization, we consider an adversarial framework where at the test time, the feature vector $\mathbf{x}$ is corrupted by additive perturbation, chosen adversarially, from the Euclidean ball of radius ε . We measure the robust generalization via the *adversarial risk* measure which is the expected test error of the model on perturbed test input. We train a random features model, using a widely used adversarial training approach, which is based on the robust empirical risk minimizer (robust-ERM estimator) [58, 90]:

$$(2.3) \quad \widehat{\boldsymbol{\theta}}^{\varepsilon} = \arg \min_{\boldsymbol{\theta} \in \mathbb{R}^N} \max_{\|\boldsymbol{\delta}_i\|_{\ell_2} \leq \varepsilon} \frac{1}{2n} \sum_{i=1}^n (y_i - \boldsymbol{\theta}^{\top} \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i)))^2.$$

where $\boldsymbol{\delta}_i$ is the norm-bounded adversarial perturbation on sample covariates $\mathbf{x}_i$ and ε is the “perceived” adversary’s power used in the training process.

Results and discussion. We study the asymptotic setting, where $N, n, d \rightarrow \infty$ with $N/d \rightarrow \psi_1$ and $n/d \rightarrow \psi_2$ for some positive constants ψ_1, ψ_2 . We derive the precise characterization of

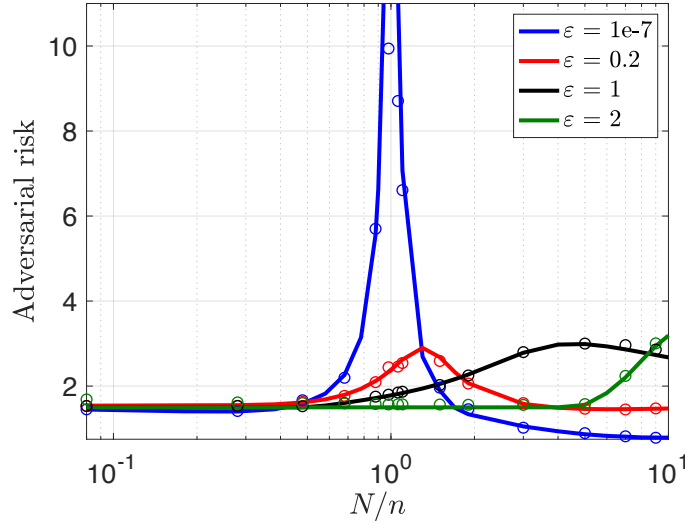


Fig 2: Adversarial risk versus overparametrization $\psi_1/\psi_2 = N/n$ for different values of adversary's power ε_0 . Solid curves are theoretical predictions and dots are results obtained based on gradient descent on the robust ERM objective. Each dot represents the average of 100 trials. The data is generated according to model (2.1), with $d = 100$, $n = 300$, $\tau^2 = 0.5$, and $\beta \in \mathbb{R}^d$ obtained by drawing a vector with i.i.d $N(0, 1)$ entries and then normalizing it to have $\|\beta\|_{\ell_2} = 1$.

the adversarial risk of the robust-ERM estimator, as an explicit function of the dimension parameters ψ_1, ψ_2 , the noise level τ^2 , and the adversarial power ε . We refer to Theorem 4.2 for the specific formulae.

Let us now discuss the behavior of the robust generalization curve under different settings. We consider the data model (2.1) and the random features regression with shifted ReLU activation:

$$\sigma(x) = \max(x, 0) - \frac{1}{\sqrt{2\pi}}.$$

The reason behind the intercept term is that since the response variable is zero mean, we consider fitting a model using zero mean features. Note that $\langle w_\ell, \mathbf{x} \rangle \sim N(0, 1)$ and for $G \sim N(0, 1)$, we have $\mathbb{E}[\sigma(G)] = \mathbb{E}[G \mathbb{I}(G > 0) - 1/\sqrt{2\pi}] = 0$.

We start by Figure 2 which shows our theoretical curve versus the overparametrization ratio $\psi_1/\psi_2 = N/n$ along with the corresponding empirical results. The solid lines depict theoretical predictions with the dots representing the empirical performance of gradient descent in learning the robust ERM for data model (2.1), with $d = 100$, $n = 300$, $\tau^2 = 0.5$. In addition, $\beta \in \mathbb{R}^d$ is generated by first drawing a d -dimensional vector with i.i.d standard normal entries and then normalizing it to have unit ℓ_2 norm. Each dot represents the average of 20 trials. As we see, even for moderate covariate dimensions (d), our theoretical curve is at excellent match with the empirical results. We note that when $\varepsilon \rightarrow 0$ (we did not set $\varepsilon = 0$ exactly for numerical stability), we are in non-adversarial regime and the robust generalization error reduces to the usual test error (blue curve). In this case, we observe the double-descent phenomena and recover the theoretical prediction of [61]. As ε grows the robust generalization curve starts behaving differently. For ε large enough ($\varepsilon = 1, 2$ in the figure), we see that overparametrization hurts robust generalization.

For a more complete picture, in Figure 3 we consider similar setting with more choices of ε and noise variance τ^2 , and also a larger range of overparametrization $\psi_2/\psi_1 = N/n$, as we fix

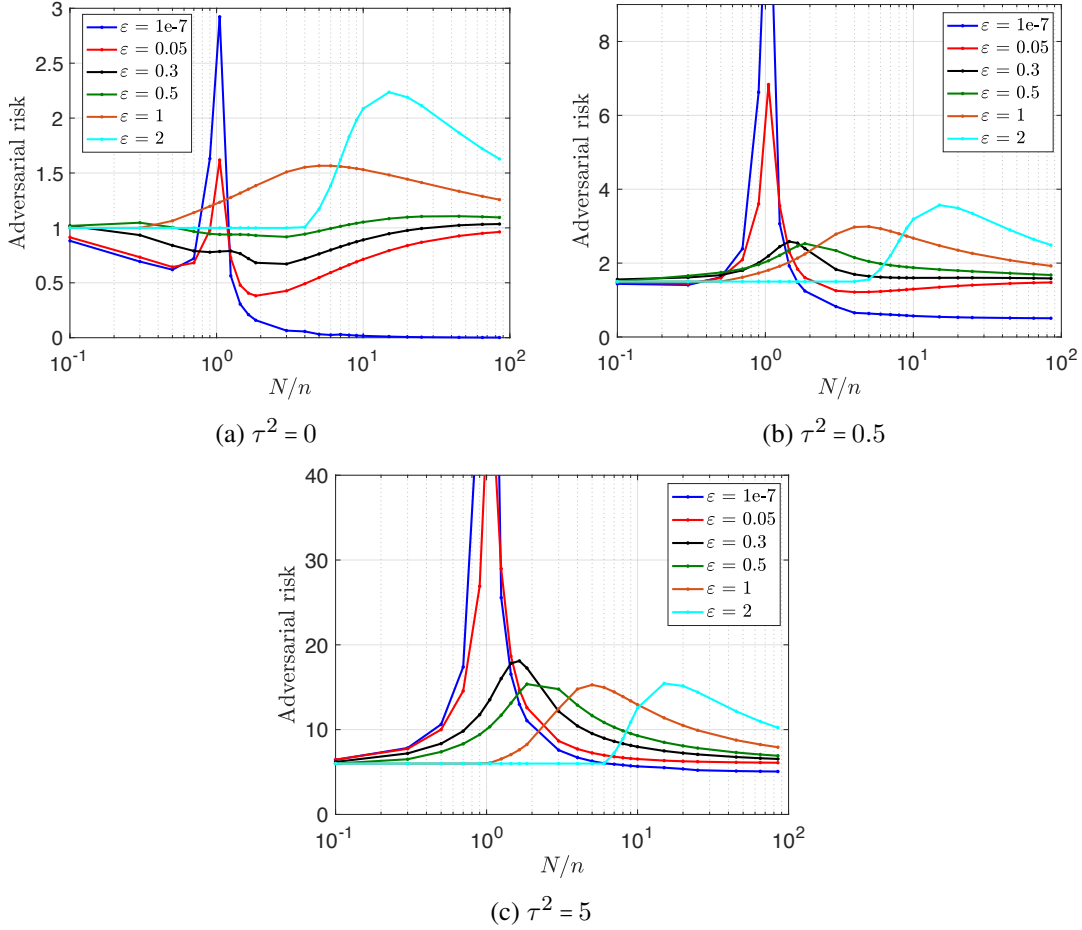


Fig 3: Theoretical prediction curves for adversarial risk of robust ERM as a function of overparametrization $\psi_1/\psi_2 = N/n$ for different values of adversary's power ε and noise variance τ^2 , with the data model (2.1). Here we fix $\|\beta\|_{\ell_2} = 1$ and $\psi_2 = 3$.

$\psi_2 = 3$. When $N/n \rightarrow 0$, we essentially have the risk of the zero estimator, which is $\|\beta\|_{\ell_2}^2 + \tau^2$. Several intriguing observations can be made from these plots:

- In the noiseless case (Figure 3a) and for $\varepsilon \leq 0.5$, the global minimizer of the adversarial risk is at a finite overparametrization ($N/n > 1$), after which the risk becomes increasing as a function of N/n (higher overparametrization hurst robust generalization). Similar behavior is observed for $\tau^2 = 0.5$ and $\varepsilon \leq 0.05$.
- In all three plots (corresponding to different SNR levels), when ε is large enough ($\varepsilon \geq 1$), the risk first goes up as overparametrization increases and after reaching its peak starts going down, but it remains above $1 + \tau^2$ which is the risk at the highly underparametrized regime ($N/n \rightarrow 0$). Therefore, somewhat surprisingly, robust ERM estimator has larger adversarial risk compared to the trivial zero estimator, for all the range of overparametrization.
- The peak of the adversarial risk occurs in the overparametrized regime; for the non-adversarial case $\varepsilon = 0$, it occurs at the interpolation threshold $N/n = 1$ and for $\varepsilon > 0$ it occurs at $N/n > 1$. The location of the peak and the value of risk at the peak vary with ε . As ε grows, the peak shifts to the right and occurs at a higher overparametrization ratio.

In Figure 4 we depict our theoretical prediction curves for the adversarial risk of the robust ERM estimator as a function of the overparametrization ratio $\psi_1/\psi_2 = N/n$ for different values

of $\psi_2 = n/d$. The right panel corresponds to $\varepsilon = 1$ (strong adversary) and as we see for different values of τ^2 and ψ_2 , the adversarial risk is first an increasing function of overparametrization ratio, until it reaches its peak (in the overparametrized regime, $N/n > 1$) and then becomes decreasing. But it never falls below its initial value at $N/n \approx 0$. The left panel corresponds to $\varepsilon = 0.1$ (weak adversary) and as we see for large ψ_2 , overparametrization clearly has a negative effect on robust generalization. For example, in Figure 4a, for $\psi_2 = 100, 1000$ the risk is an increasing function of ψ_1/ψ_2 over the entire range. Also in Figure 4c, for $\psi_2 \geq 10$ the global minimum of the adversarial risk is achieved in the underparametrized regime ($N/n < 1$).

3. Related Work Several recent works have focused on the robustness of overparametrized models. On the one hand, [9] shows that in order *interpolate* the training data smoothly, the Lipschitz parameter of the resulting model should be at least of order $\sqrt{\frac{nd}{N}}$. This applies to data distributions that satisfy a property called isometry—e.g. when the data covariates $\mathbf{x}_i$ are distributed on the unit sphere. For such data distributions, worst-case perturbations are meaningful only if their ℓ_2 norm is upper-bounded by $\frac{\epsilon}{\sqrt{d}}$. Otherwise, if the size of the perturbation can be allowed to be much larger than $O(\frac{1}{\sqrt{d}})$, it can be shown that the robust generalization error approaches one for any model—see [75, 29, 59, 60]. Putting the above two results together, we can conclude that, in order to interpolate smoothly, while guaranteeing robustness to norm-bounded perturbations, it is necessary that the ratio N/n is bounded away from zero. This is indeed the regime studied in our paper. However, in this regime, it is not clear why interpolation to training data is beneficial for obtaining robust models. In fact, to obtain robust models, one may have to trade off the performance on the original data points (i.e. interpolation to training data) with the performance on the points in a ball around each data point (i.e. extrapolation to adversarial examples). In other words, we may have underparametrized models that do not fit the training data perfectly, but have a small Lipschitz constant. Indeed, this can be implied from the main messages of our paper.

On the other hand, the works in [46] and [23] have studied the performance of high-dimensional *linear* models and showed that the robust generalization error of adversarially-trained models becomes worse as the models become more overparametrized. In particular, [23] provably shows that avoiding interpolation (and using underparametrized models) improves the robust generalization error in both linear regression and classification—which leads to the first theoretical result on robust overfitting. There are a few reasons on why we might prefer to study non-linear models (such as the random features model) compared to linear models [61]: First of all, for linear models, we know that the best (standard) generalization error is attained when the model is highly underparametrized. Second, the number of parameters in a linear model is tied to the covariates dimension d and hence the effects of overparametrization cannot be isolated from the effects of the ambient dimensions. Third, a hypothesis put forward in [32] is that the origins and ubiquity of adversarial examples is due to the (approximately) linear behavior of a model over large regions of the input space. Shallow linear models are not able to become constant near training points while also assigning different outputs to different training points. However, the setting of random features is significantly different since this class can express any function to an arbitrary degree of accuracy so long as it has enough number of random features [69].

In another recent work [94], an extensive study on the robustness of wide neural networks with respect to norm-bounded perturbations is provided. By defining and analyzing a new metric, called perturbation stability, it is shown that while the (standard) generalization error is improved on wider models, the perturbation stability often worsens, leading to a potential decrease in the overall model robustness. These empirical findings are aligned with the messages of our paper.

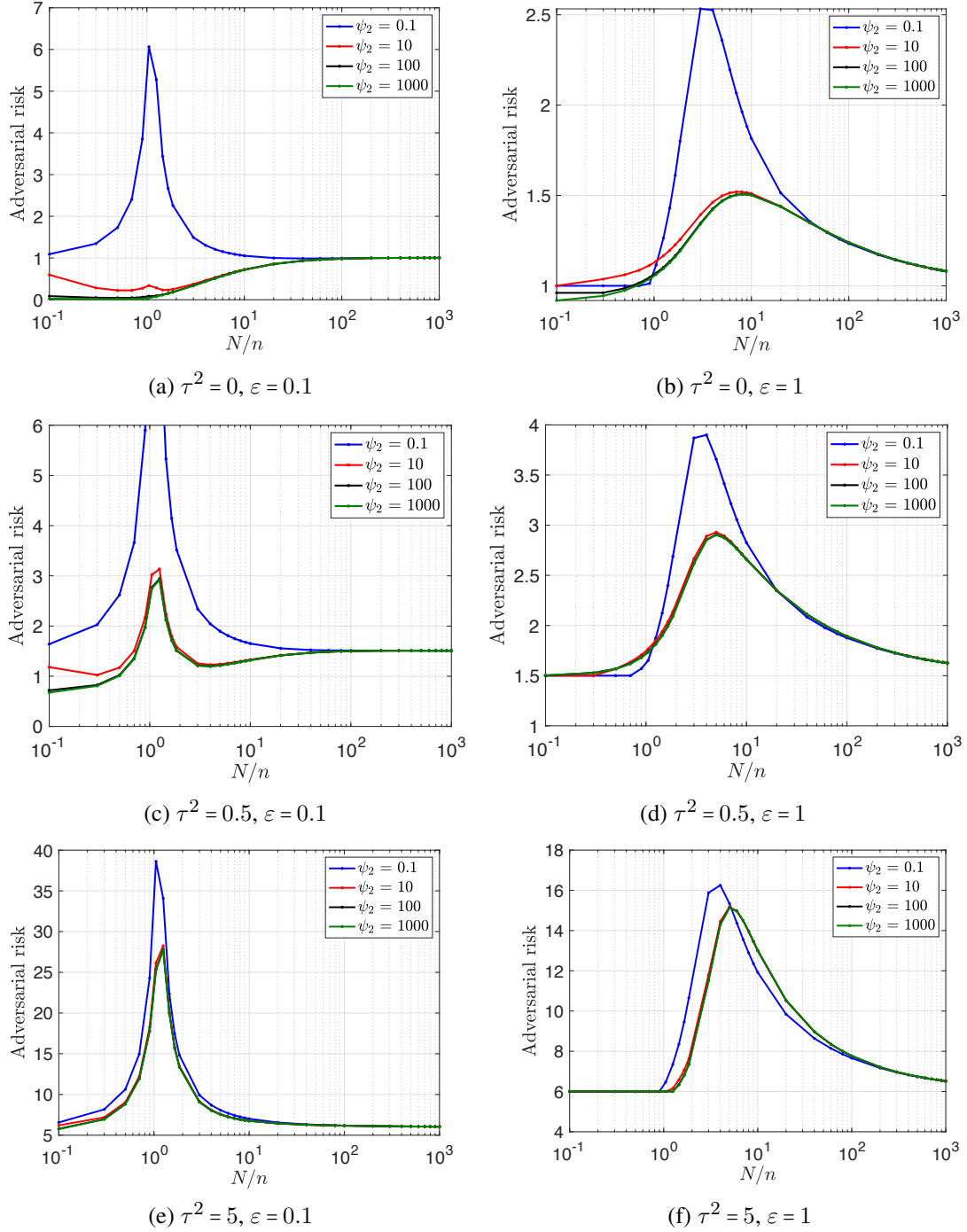


Fig 4: Theoretical prediction curves for adversarial risk of robust ERM as a function of overparametrization $\psi_1/\psi_2 = N/n$ for different values of $\psi_2 = n/d$. Each plot corresponds to a specific value of adversary's power ε and noise variance τ^2 .

A somehow different line of work [96] studies the sample complexity of the robust interpolation problem where the goal is to interpolate (noisy) training data by a Lipschitz function, under generic covariate distribution (beyond isoperimetry distributions). This work measures the (non)robustness of a model by its Lipschitz constant and similar to [22] establish a lower

bound on Lipschitzness which is increasing with respect to the overfitting level. This result can be rephrased as an adverse effect of overparametrization (thorough overfitting) on robustness. While these work study the effect of overparametrization on robustness via memorization/interpolation, we will take a direct approach to study the effect of overparametrization on ‘adversarially trained’ models.

Several works have shown a non-trivial tradeoff between the robust generalization error and the standard generalization error for parametric models [90, 83, 67, 97, 46, 21]. It has also been shown that using more data can improve this tradeoff [11, 62, 67, 74, 17, 95, 65, 34, 70]. Again, these findings are aligned with the messages of our paper as more data can mean less overparametrization.

This paper provides, for the first time, an analysis for the adversarially-trained random features model in the high-dimensional regime. For linear models, this analysis has been carried out in [46] for the regression setting, and later on in [45, 85] for the classification setting. A key ingredient of the analysis in these papers, as well as our paper, is a powerful extension of a classical Gaussian process inequality [33], known as the Convex Gaussian Minimax Theorem, developed in [88] and further extended in [87, 16]. Another key ingredient of our analysis is the Gaussian Equivalence Property for the random features model which was proposed and studied in [63] for maximum-margin linear classifiers in the overparametrized regime, as well as [37, 1, 63, 61, 28, 19, 39, 31, 30] for the linear Gaussian model under other settings. In particular, a part of our analysis, that establishes equivalence with the so-called noisy linear model in the adversarial setting, is heavily based on the machinery which was elegantly developed in [39] for the random features model. This machinery is itself based on the Lindeberg principle [55] and the leave-one-out technique developed in [27, 1].

We conclude this section by a broad comparison between adversarial setting and the literature of robust statistics.

Comparison with robust statistics. This area traditionally considers a setting where perturbations are made to the *training data*; a small fraction of data samples are grossly corrupted and the goal is to find estimators that are robust against outliers (via measures like influence function, breakdown point, and change of variance, etc). In the adversarial training paradigm, one considers the so-called test-time adversarial setting, in which the training data is uncorrupted (say $(x_i, y_i) \sim \mathbb{P}$ for some distribution $\mathbb{P}$). However, the adversary can perturb *each test data*. In other words, the test data (x, y) is drawn from $\mathbb{P}$ and then x is perturbed by the adversary ($x \rightarrow \tilde{x}$ with $\|x - \tilde{x}\|_{\ell_2} \leq \varepsilon$). The goal of adversarial training is to develop a model that can still predict the response y from the perturbed feature $\tilde{x}$. With this view, adversarially robust models are basically those that have good generalization and are also smooth enough so that they do not change much on small neighborhoods (of radius at most ε).

That said, another line of work (see e.g. [48]) considers a different adversarial setup in which an attacker can observe and modify all training data samples in an adversarial manner so as to maximize the estimation error caused by his attack. This work introduces the notion of adversarial influence function (AIF) to quantify the sensitivity of estimators to such adversarial attacks, and further derive the optimal estimator, among a certain class of estimator, that minimizes AIF. Related to this setting, there is also a line of work based on the Median of Means approach, see e.g. [40, 18]), which concerns a data poisoning/ data contamination adversarial setting. In data poisoning, the adversary can pick a (small) fraction of the *training data* and alter it in a way that it hurts the training process, and ultimately the generalization performance. However, in this paper we consider a different type of adversarial act which has to do with adversarial perturbation (in a small ball) of the input data point at the *test time*.

4. Main results Recall the data distribution given in (2.1). Given n i.i.d pairs $(\mathbf{x}_i, y_i)$ drawn from this distribution, we fit a random features model, defined as the function class (2.2), with the shifted ReLU activation:

$$(4.1) \quad \sigma(x) = \max(x, 0) - \frac{1}{\sqrt{2\pi}}.$$

We consider sequences of parameters (N, n, d) that diverge proportionally to each other and sometimes, we index such sequences by d , with $N = N(d)$ and $n = n(d)$ functions of d .

Assumption 1 (*Asymptotic setting.*)

(a) Defining $\psi_{1,d} = N/d$ and $\psi_{2,d} = n/d$, we assume that the following limits exist:

$$\lim_{d \rightarrow \infty} \psi_{1,d} = \psi_1, \quad \lim_{d \rightarrow \infty} \psi_{2,d} = \psi_2,$$

for some positive finite constants ψ_1 and ψ_2 .

(b) We assume that the ℓ_2 norm of the signal β converges, as $d \rightarrow \infty$. For the sake of normalization and without loss of generality, we assume $\lim_{d \rightarrow \infty} \|\beta\|_{\ell_2} = 1$.

Recall that in the data model (2.1), $\mathbf{x}_i \sim_{iid} \mathcal{N}(0, \mathbf{I}_d)$ and so its distribution is rotation-invariant. Likewise, in the random features model (2.2), the rows $\mathbf{w}_\ell$ are chosen uniformly at random from the unit sphere, and so has a rotation-invariant distribution. In our adversarial setting, we also focus on norm-bounded perturbations which is again a rotation-invariant constraint. Using these properties, it is easy to see that the adversarial risk will be invariant if we rotate the model β in (2.1) and hence only depends on $\|\beta\|_{\ell_2}$. This justifies Assumption 1(b) made above.

To study robust generalization of the estimated models, we consider an adversarial framework with norm bounded perturbations. This can be formulated as a game between the learner and the adversary. Given access to a training dataset consisting of n i.i.d. pairs $(\mathbf{x}_i, y_i)$ generated from (2.1), the learner chooses a model θ from the class of random features model $\mathcal{F}_{\text{RF}}(\mathbf{W})$ (2.2). Adversarial perturbations occur at the test time. After observing the learner's model, for every test point $\mathbf{x}$, the adversary perturbs it to $\mathbf{x} + \delta$ where δ is chosen from the Euclidean ball of radius ε . Note that the choice of δ can in general depend on $\mathbf{x}$ and the learner's model. The robust generalization of the learner's model is quantified via a measure called *adversarial risk*, which is the expected prediction loss of the model on an adversarially corrupted test data point according to the described attack model.

Definition 4.1 (*Adversarial risk.*) For a predictive model f and a loss of choice $\ell : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}_{\geq 0}$, the adversarial risk of model f is defined as:

$$\text{AR}(f) := \mathbb{E} \left[\max_{\|\delta\|_{\ell_2} \leq \varepsilon_{\text{test}}} \ell(f(\mathbf{x} + \delta), y) \right],$$

where the expectation is with respect to randomness of $(\mathbf{x}, y)$.

In particular, for a random features model $\theta = (\theta_1, \dots, \theta_N)^\top$ from $\mathcal{F}_{\text{RF}}(\mathbf{W})$, defined in (2.2), and with the choice of squared loss, the adversarial risk of θ becomes:

$$(4.2) \quad \text{AR}(\theta) := \mathbb{E} \left[\max_{\|\delta\|_{\ell_2} \leq \varepsilon_{\text{test}}} (y - \theta^\top \sigma(\mathbf{W}(\mathbf{x} + \delta)))^2 \right].$$

Norm bounded adversarial attack models are widely used in the literature, motivated by a plethora of safety-critical applications in machine learning, computer vision, natural language processing, medical imaging, and robotics. A popular approach to adversarial training is by considering the following robust empirical risk minimization (robust-ERM) problem [58, 90]:

$$(4.3) \quad \widehat{\theta}^\varepsilon = \arg \min_{\theta \in \mathbb{R}^N} \max_{\|\delta_i\|_{\ell_2} \leq \varepsilon} \frac{1}{2n} \sum_{i=1}^n (y_i - \theta^\top \sigma(\mathbf{W}(\mathbf{x}_i + \delta_i)))^2.$$

Here, $\varepsilon_{\text{test}}$ is a measure of the adversary's power and ε is the “perceived” adversary's power used by the algorithm. Our theory allows for ε to be different from $\varepsilon_{\text{test}}$; cf Theorem 4.2. In our numerical experiments in Section 2 we consider $\varepsilon = \varepsilon_{\text{test}}$ to focus on other relevant quantities, namely ψ_1, ψ_2 on adversarial risk.

Note that the above objective is the empirical surrogate of the adversarial risk (4.2), where the expectation is replaced by the empirical average over the training samples. This minimax approach can also be viewed as an implicit smoothing that tries to fit the same response y to all the covariate vectors in the ε -neighborhood of $\mathbf{x}$ simultaneously.

Our main result in this paper is a precise characterization of the adversarial risk of the robust ERM model (4.3) under the asymptotic regime described in Assumption 1. Before stating our result, we introduce another piece of notation.

For $\psi_1 \in (0, \infty)$, we define function $S(\cdot; \psi_1) : \mathbb{R}_{<0} \rightarrow \mathbb{R}_{<0}$:

$$(4.4) \quad S(z; \psi_1) = \frac{1 - \psi_1 - z - \sqrt{(1 - \psi_1 - z)^2 - 4\psi_1 z}}{-2\psi_1 z}.$$

One may recognize that $S(z; \psi_1)$ is the Stieltjes transform of the Marchenko-Pastur distribution. We refer to Lemma F.2 (Appendix F) for more details.

We are now ready to state our main result.

Theorem 4.2 *Let n i.i.d pairs $(\mathbf{x}_i, y_i)$ be drawn from the data model (2.1) and let $\widehat{\boldsymbol{\theta}}^\varepsilon$ be the robust ERM fit (4.3) to this data using the class of random features models $\mathcal{F}_{\text{RF}}(\mathbf{W})$, given by (2.2) with the shifted ReLU activation. Consider the asymptotic regime, described in Assumption 1. With function $S(\cdot; \psi_1)$ given by (4.4), define*

$$\sigma^2 = \tau^2 + 1 - \psi_1 \left(1 + \left(1 - \frac{2}{\pi} \right) S\left(\frac{2}{\pi} - 1; \psi_1\right) \right).$$

(a) *For $\varepsilon > 0$, the following convex-concave minimax scalar optimization has a unique solution $(\alpha_*, \tau_{g*}, \beta_*, \gamma_*, \tau_{q*})$:*

$$(4.5) \quad \max_{0 \leq \beta, \gamma, \tau_q} \min_{0 \leq \alpha, \tau_g} \mathcal{R}(\alpha, \tau_g, \beta, \gamma, \tau_q),$$

where

$$\begin{aligned} \mathcal{R}(\alpha, \tau_g, \beta, \gamma, \tau_q) := & \frac{\tau_q}{2\alpha} (\tau^2 + 1 - \sigma^2) - \frac{\alpha\tau_q}{2} + \frac{\beta\tau_g}{2} \psi_2 + \frac{\beta}{2(\tau_g + \beta)} (\sigma^2 + \alpha^2) \\ & + \mathbf{1}_{\left\{ \frac{\gamma(\tau_g + \beta)}{\varepsilon\beta\sqrt{\alpha^2 + \sigma^2}} > \sqrt{\frac{2}{\pi}} \right\}} \frac{\beta^2(\alpha^2 + \sigma^2)}{2\tau_g(\tau_g + \beta)} \left(\operatorname{erf}\left(\frac{\nu^*}{\sqrt{2}}\right) - \frac{\gamma(\tau_g + \beta)}{\varepsilon\beta\sqrt{\alpha^2 + \sigma^2}} \nu^* \right) \\ & - \frac{\alpha}{\tau_q} \sup_{0 \leq \lambda < 1} \left[\frac{\lambda\psi_1}{2} \left\{ \frac{\tau_q^2}{\alpha^2} + \beta^2 + \left(\frac{\tau_q^2}{\alpha^2} \left(1 - \frac{2}{\pi} \lambda \right) + \frac{2}{\pi} (1 - \lambda) \beta^2 \right) S\left(\frac{2}{\pi} \lambda - 1; \psi_1\right) \right\} - \frac{\lambda}{2(1 - \lambda)} \gamma^2 \right]. \end{aligned}$$

Here, ν^* is the unique solution to

$$\frac{\gamma(\tau_g + \beta)}{\varepsilon\beta\sqrt{\alpha^2 + \sigma^2}} - \frac{\beta}{\tau_g} \nu - \nu \cdot \operatorname{erf}\left(\frac{\nu}{\sqrt{2}}\right) - \sqrt{\frac{2}{\pi}} e^{-\frac{\nu^2}{2}} = 0.$$

(b) *The adversarial risk of the robust ERM $\widehat{\boldsymbol{\theta}}^\varepsilon$ converges in probability*

$$(4.6) \quad \text{AR}(\widehat{\boldsymbol{\theta}}^\varepsilon) \xrightarrow{\mathcal{P}} \left[1 + \left(\frac{\varepsilon_{\text{test}} \beta_* \nu_*}{\varepsilon \tau_{g*}} \right)^2 + 2 \sqrt{\frac{2}{\pi}} \frac{\varepsilon_{\text{test}} \beta_* \nu_*}{\varepsilon \tau_{g*}} \right] (\alpha_*^2 + \sigma^2).$$

Here, the probabilistic statement is with respect to the randomness in both the training data $\{(\mathbf{x}_i, y_i)\}_{i=1}^n$ and the random features $\mathbf{W}$.

We note that the robust ERM estimator is a *random* and rather complicated high-dimensional function of the training data. However, in the asymptotic regime where $N, n, d \rightarrow \infty$ at the same order, the adversarial risk of the robust ERM concentrates and the above theorem provides an exact characterization of its limit as a *deterministic* formula. The derived formula is based on a five dimensional convex-concave mini-max optimization problem and its optimal solution can be easily derived via a simple low-dimensional gradient descent rather quickly and accurately. Alternatively, one can form a system of equations by writing the KKT stationary conditions corresponding to (4.5). The adversarial risk prediction can then be written in terms of the fixed point of this system of deterministic equations.

Let us re-emphasize the contribution of theorem 4.2. Albeit its involved form, it describes the behavior of a *high-dimensional random* problem in terms of a *deterministic* optimization with a handful number of scalar variables. This theme of result is similar to state evolution equation for approximate message passing algorithms [24], density evolution for LDPC codes [71, Chapter 4], and characterizing the trajectory of SGD for training neural networks in terms of partial differential equations [78, 44].

Remark 4.1 (Solving Optimization (4.5)) *We find the solution to this optimization by solving for the first-order optimality conditions (stationary equations). We set the (sub)gradient with respect to the variables to zero since it is non-smooth, which results in a system of nonlinear equations with seven variables, namely $\alpha, \tau_g, \beta, \gamma, \tau_q, \lambda, \nu^*$. In the numerical experiments, we use `fsolve` command in `Matlab` to solve this system of equations, which is based on the trust-region algorithm.*

5. Discussion By virtue of Theorem 4.2 we characterize the adversarial risk of the robust ERM in term of the solution of the deterministic optimization problem (4.5). Given that it does not admit a closed-form solution in general, and is rather involved, in this section we discuss some of the applications of this theorem including optimal choice of ε during training, trend of adversarial risk with respect to different quantities, and implications for non-adversarial setting.

5.1. Optimal ε for robust ERM estimator: An interesting application of our theory is to derive the optimal ε (perceived adversary’s perturbation level) in the robust ERM, while fixing the adversary’s (actual) perturbation level on test inputs to $\varepsilon_{\text{test}}$. The optimal ε here refers to the value which minimized the adversarial risk. An intriguing observation is that the optimal ε is different than $\varepsilon_{\text{test}}$ in general, and depends on ψ_1, ψ_2 in a non-trivial way (There is no universal solution, which underscores the significance of possessing a precise theory that comprehends the impact of adversarial training, which constitutes the principal objective of the present work.)

In Figure 5a, we fix $\psi_2 = 3$, $\tau = \sqrt{0.5}$, $\varepsilon_{\text{test}} = 0.3$, and plot the adversarial risk of $\widehat{\theta}^\varepsilon$ as we vary ε for different values of ψ_1 . As we see the optimal value of ε (resulting in minimum risk) changes with ψ_1 , it is in general different from the test adversary’s perturbation $\varepsilon_{\text{test}}$. In addition, the optimal ε increases with ψ_1 . In Figure 5b, we plot similar curves, fixing $\varepsilon_{\text{test}} = 0$. In this case, the adversarial risk reduces to the notion of standard risk defined as

$$(5.1) \quad \text{SR}(\theta) := \mathbb{E}[(y - \theta^\top \sigma(Wx))^2].$$

As we see from the plots, even though $\varepsilon_{\text{test}} = 0$, adversarial training can help as minimum risk is achieved at positive ε . The reason is that adversarial training acts as a regularization (It becomes clearer after derivation (6.11), where adversarial training aims to find solution with small $\|J\theta\|_{\ell_2}$.) In particular, we observe that

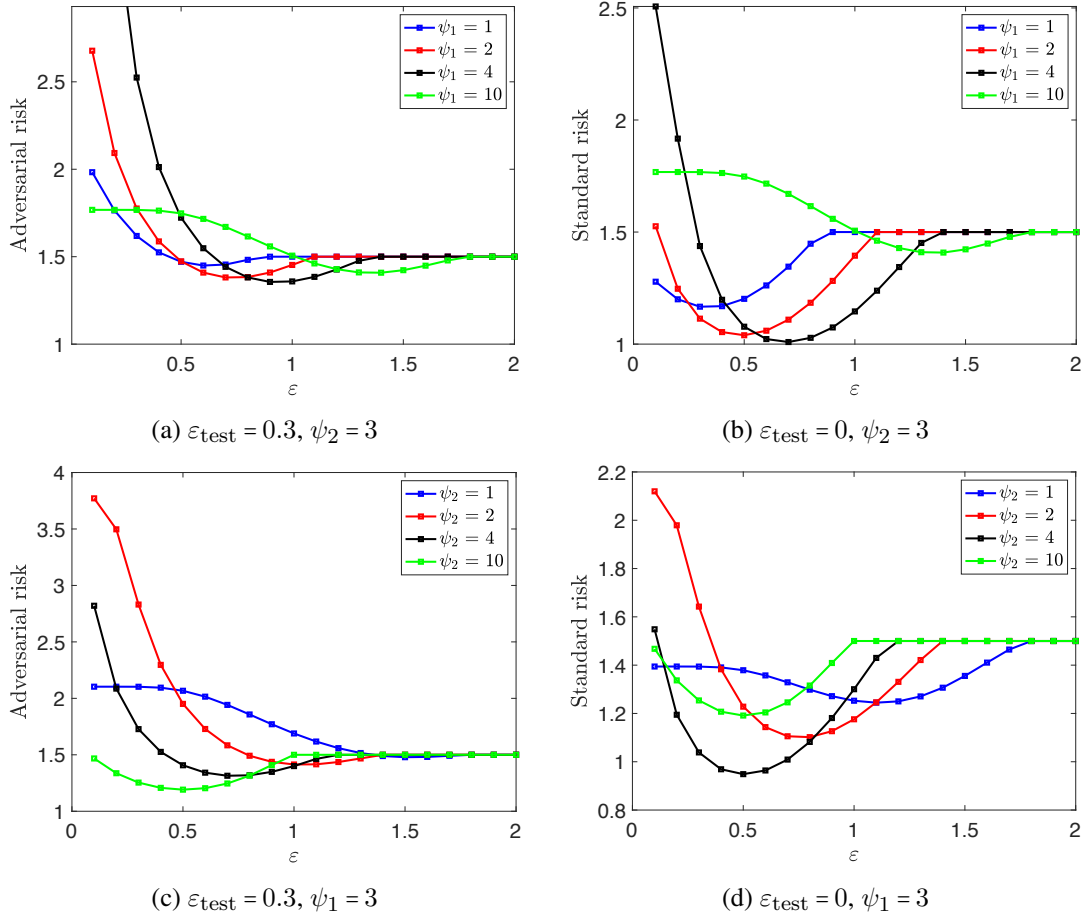


Fig 5: Behavior of adversarial/standard risk as we vary ε , the “perceived” adversary’s power used in the adversarial training. In (a), (b), $\psi_2 = 3$ is fixed and in (c), (d), we fix $\psi_1 = 3$. Also, (b), (d) correspond to $\varepsilon_{\text{test}} = 0$, and so there is no perturbation at the test time. In these cases, adversarial risk reduces to the standard risk. In these experiments, we set $\tau^2 = 0.5$, $\|\beta\|_{\ell_2} = 1$.

- The optimal ε is always greater than or equal to $\varepsilon_{\text{test}}$, the true test perturbation magnitude. This ‘additional’ regularization helps with minimizing the adversarial risk.
- At higher overparametrization measured by $\psi_1/\psi_2 = N/n$ the benefit of this regularization becomes stronger. This is also evident from Figure 5c, where with fixed ψ_2 . As we increase ψ_1 , the optimal ε also increases. Likewise, in Figure 5d, with fixed ψ_1 , the optimal ε increases as ψ_2 decreases.

In summary, our theory allows to understand when adversarial training is beneficial and what is the optimal value of ε to use in training (depending on $\psi_1, \psi_2, \varepsilon_{\text{test}}$ and τ .)

5.2. Dependence on $\varepsilon, \psi_1, \psi_2$: We discern the following trends in the analytical curves for adversarial risk which are derived based on Theorem 4.2 (Equation (4.6)).

- From Figure 5, fixing $\varepsilon_{\text{test}}, \psi_1$ and ψ_2 , the adversarial risk is first decreasing in ε until it gets to its minimum, after which it becomes increasing in ε indicating that the regularization effect from adversarial training is larger than it should be, and then eventually it levels

at $\sigma^2 + \|\theta_0\|^2$ (the risk of model $\theta = \mathbf{0}$, which is the ERM when the regularization is very strong).

- For fixed ψ_2 , if SNR is large enough and ε is small enough, we observe a double descent behavior in the adversarial risk (see Figure 3a, $\varepsilon = 1e-7$ and $\varepsilon = 0.05$ and Figure 4c).
- Fixing $\varepsilon_{\text{test}}$ and ψ_1 , the adversarial risk becomes decreasing in ψ_2 after some point. In the next subsection, we characterize the adversarial risk in non-adversarial setting, by which we see this trend for $\psi_2 > \psi_1$.

5.3. Non-adversarial training: An important special case of the result is when $\varepsilon = \varepsilon_{\text{test}} = 0$. In words, there is no adversarial-training and also there is no adversary during the test time. Theorem 4.2 allows us to characterize the standard risk (generalization error) of the ERM estimator.

We focus on the underparameterized regime ($n > N$ or equivalently $\psi_2 > \psi_1$), since otherwise at $\varepsilon = 0$ the problem is underdetermined. In this case the objective $\mathcal{R}$ in (4.5) significantly simplifies and we can indeed obtain a closed-form expression for the risk.

Proposition 5.1 *Let n i.i.d pairs $(\mathbf{x}_i, y_i)$ be drawn from the data model (2.1) and let $\hat{\theta}$ be the ERM (4.3) fit to this data using the class of random features models $\mathcal{F}_{\text{RF}}(\mathbf{W})$, given by*

$$\hat{\theta} = \arg \min_{\theta \in \mathbb{R}^N} \frac{1}{2n} \sum_{i=1}^n (y_i - \theta^\top \sigma(\mathbf{W} \mathbf{x}_i))^2,$$

with $\sigma(\cdot)$ the shifted ReLU activation. Consider the asymptotic regime, described in Assumption 1. With function $S(\cdot; \psi_1)$ given by (4.4), define

$$\sigma^2 = \tau^2 + 1 - \psi_1 \left(1 + \left(1 - \frac{2}{\pi} \right) S\left(\frac{2}{\pi} - 1; \psi_1\right) \right).$$

Then, the standard risk of the ERM $\hat{\theta}$ converges in probability to

$$(5.2) \quad \text{SR}(\hat{\theta}) \xrightarrow{\mathcal{P}} \sigma^2 \left(\frac{\psi_2}{\psi_2 - \psi_1} \right).$$

We refer to Section E.4 for the proof of this proposition.

We next use the result of Proposition 5.1 to discuss the role of ψ_1 and ψ_2 on the risk:

- Recall that σ only depends on τ and ψ_1 . Fixing ψ_1 the dependence of risk on ψ_2 is of form $\psi_2/(\psi_2 - \psi_1)$. This is decreasing in $\psi_2 = n/d$. For example if d, N are fixed, and we increase the sample size n , the risk goes down which is expected.
- Dependence on ψ_1 is more involved as the term σ^2 also depends on ψ_1 through the Stieltjes transform S . In Figure 6 we plot the risk versus ψ_1 for different values of ψ_2 . As we see up to some threshold, it is decreasing in ψ_1 but after that it becomes increasing. This is expected because for example fixing n, d (and so ψ_2), as we increase N (and so ψ_1), first the risk goes down because the model becomes richer to capture the data generative model, but after some point it has a reverse effect, because we need to estimate larger number of parameters N , from fixed sample size n , while this excess model complexity is not needed. As we see in the plots, this threshold is increasing with ψ_2 .
- It is also worth comparing the standard risk of random features model with that of linear models. For $\psi_2 \geq 1$, using the result of [37, Proposition 2], the risk of ridgeless least squares is given by $\tau^2 \psi_2 / (\psi_2 - 1)$. This is similar to our characterization (5.2), where the noise variance τ^2 is replaced with the effective noise variance σ^2 , and ψ_2 is replaced by $\psi_2 / \psi_1 = n/N$. (Note that the number of parameters to be learnt in the linear model is d , while in the random features model is N . So the sample size to parameter size ratio in the linear regression is ψ_2 , while for the random features model it is ψ_2 / ψ_1 .)

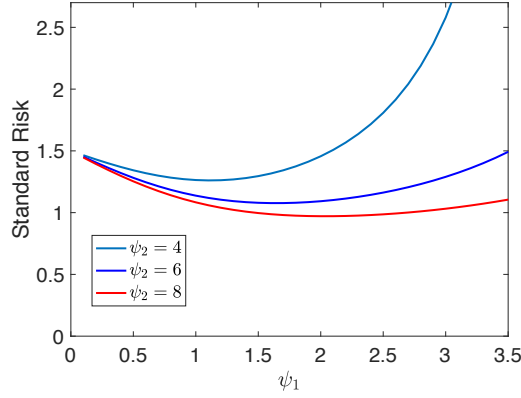


Fig 6: Behavior of the standard risk of the ERM estimator $\widehat{\theta}$ versus ψ_1 for different values of ψ_2 . As observed, the risk is first decreasing in ψ_1 , up to some threshold depending on ψ_2 , after which it becomes increasing. This threshold is increasing with ψ_2 .

6. Architecture of the proof This section introduces the key steps underlying the proof of our main result, Theorem 4.2. Our analysis is intricate and consists of a host of novel ideas which could be of separate interest. Here we discuss the major steps, along with an overview of the techniques and intermediate results.

Define the loss $\mathcal{L}(\theta)$ given by

$$(6.1) \quad \mathcal{L}(\theta) := \max_{\|\delta_i\|_{\ell_2} \leq \varepsilon} \frac{1}{2n} \sum_{i=1}^n (y_i - \theta^\top \sigma(\mathbf{W}(\mathbf{x}_i + \delta_i)))^2 + \frac{\zeta}{2} \theta^\top \Omega \theta,$$

where $\Omega := \mathbf{I} + \frac{\sqrt{\log(d)}}{d} \mathbf{1}\mathbf{1}^\top$. Here, $\mathbf{1} = (1, 1, \dots, 1) \in \mathbb{R}^N$ and $\zeta > 0$ is an arbitrary small but fixed constant. By definition, the robust ERM estimator (4.3) is the minimizer of $\mathcal{L}(\theta)$ for $\zeta = 0$.

The regularization $\frac{\zeta}{2} (\|\theta\|_{\ell_2}^2 + \frac{\sqrt{\log(d)}}{d} \langle \mathbf{1}, \theta \rangle^2)$ in the loss $\mathcal{L}(\theta)$ is added for technical reasons. In our analysis, we let $\zeta \rightarrow 0$, after letting $d \rightarrow \infty$ to characterize the adversarial risk of the robust ERM $\widehat{\theta}^\varepsilon$. We refer to Section A in the supplementary for the justification of this step.

Before we outline the main steps of the proof, we note that since the rows of the matrix $\mathbf{W} \in \mathbb{R}^{N \times d}$ are generated i.i.d. according to $\mathbf{w}_\ell \sim \text{Unif}(\mathbb{S}^{d-1})$, then the matrix norm of $\mathbf{W}$ is bounded with high probability and the rows of $\mathbf{W}$ are almost orthogonal. More precisely, we define the event

$$(6.2) \quad \mathcal{E}_{\mathbf{W}} := \left\{ \|\mathbf{W}\| \leq \sqrt{\psi_{1,d}} + C, \quad |\mathbf{w}_\ell^\top \mathbf{w}_k| \leq \log(d)/\sqrt{d} \quad \forall \ell \neq k \right\},$$

for a large enough constant C (note that $N/d := \psi_{1,d}$ – see Assumption 1). Using well-known results on the norm of random matrices (see e.g. [91, Theorem 5.39]) as well as Hoeffding's inequality we have $\mathbb{P}(\mathcal{E}_{\mathbf{W}}) \geq 1 - c \exp(-\log^2(d)/c)$ for some constant $c > 0$. In the following, our statements are proven conditioned on the event $\mathcal{E}_{\mathbf{W}}$ which holds with high probability.

Notation. We need to define a few pieces of notation which will be used in the following. We use $O_d(\cdot)$, $o_d(\cdot)$ to denote the standard big-O and little-o notation, where we stress the asymptotic variable d . Likewise, we use $O_{d,\mathbb{P}}$ and $o_{d,\mathbb{P}}$ to indicate asymptotic behavior in probability. Specifically, $f(d) = O_{d,\mathbb{P}}(g(d))$ if for any $\varepsilon > 0$, there exists $C_\varepsilon > 0$ and large enough d_ε such that $\mathbb{P}(|f(d)/g(d)| > C_\varepsilon) \leq \varepsilon$, for all $d \geq d_\varepsilon$. Similarly, $f(d) = o_{d,\mathbb{P}}(g(d))$ if $f(d)/g(d)$ converges to zero in probability. We write $f(d) \approx g(d)$ as $d \rightarrow \infty$, when $f(d) -$

$g(d) \rightarrow 0$, in probability. Note that we consider the asymptotic regime where n, d, N grow at the same scale, ($\lim N/d \rightarrow \psi_1$ and $\lim n/d \rightarrow \psi_2$ for some positive constants ψ_1 and ψ_2), the expression $d \rightarrow \infty$ implies that $n, N \rightarrow \infty$, as well.

For a matrix $\mathbf{A}$, we denote by $\|\mathbf{A}\|$ its operator norm, $\|\mathbf{A}\|_F = (\sum_{ij} A_{ij}^2)^{1/2}$ the Frobenius norm of $\mathbf{A}$. For an integer n , we use the shorthand $[n] = \{1, \dots, n\}$.

Finally, the indicator function is denoted by $\mathbb{I}(\cdot)$ – i.e., $\mathbb{I}(A) = 1$ only if the event A holds true, and otherwise $\mathbb{I}(A) = 0$.

6.1. An asymptotically-exact closed form for adversarial examples We start by simplifying the loss $\mathcal{L}(\boldsymbol{\theta})$. If the activation function σ was linear, then finding the worst-case perturbations $\boldsymbol{\delta}_i$ (maximizers in the definition of loss $\mathcal{L}(\boldsymbol{\theta})$) amounts to a trust-region subproblem that can be solved in closed form—see [46]. A major challenge here is the nonlinearity of σ . In the first step we use the specific form of the activation to derive an asymptotically equivalent but simpler form of $\mathcal{L}(\boldsymbol{\theta})$.

Note that $\sigma(z) = \max(z, 0) - 1/\sqrt{2\pi}$ is linear on the positive z and constant on the negative z . Also by the constraint on perturbation we have $\|\langle \mathbf{w}, \boldsymbol{\delta} \rangle\| \leq \|\mathbf{w}\|_{\ell_2} \|\boldsymbol{\delta}\|_{\ell_2} \leq \varepsilon$. Therefore, if $\|\langle \mathbf{w}, \mathbf{x} \rangle\| \geq \varepsilon$, then $\langle \mathbf{w}, \mathbf{x} \rangle$ and the perturbed form $\langle \mathbf{w}, \mathbf{x} + \boldsymbol{\delta} \rangle$ share the same sign. In this case, the worst case $\boldsymbol{\delta}$ can be solved exactly. One can also use the randomness in $\mathbf{x}$ to bound the number of rows of $\mathbf{W}$ for which $|\langle \mathbf{w}, \mathbf{x} \rangle| < \gamma$, where γ is some small constant, and show that the contribution of these terms in the loss is asymptotically negligible. This argument is formalized in the next proposition. All the proofs of the statements in this section are relegated to Appendix B.

Proposition 6.1 Assume $(\mathbf{x}, y)$ generated according to (2.1). Further define

$$(6.3) \quad \mathcal{C}_\theta := \{\boldsymbol{\theta} \in \mathbb{R}^N : \|\boldsymbol{\theta}\|_{\ell_\infty} \leq C_0 \sqrt{\log(d)/d}, \|\boldsymbol{\theta}\|_{\ell_2} \leq C_0\},$$

for an arbitrary but fixed constant $C_0 > 0$. Then, we have

$$(6.4) \quad \max_{\|\boldsymbol{\delta}\|_{\ell_2} \leq \varepsilon} |y - \boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x} + \boldsymbol{\delta}))| = |y - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x})| + \varepsilon \|\mathbf{W}^\top \text{diag}(\mathbb{I}(\mathbf{W}\mathbf{x} > 0)) \boldsymbol{\theta}\|_{\ell_2} + O_{d, \mathbb{P}}\left(\frac{\log(d)}{d^{1/6}}\right),$$

uniformly over $\boldsymbol{\theta} \in \mathcal{C}_\theta$. Here, the probability bound is with respect to the randomness in $\mathbf{x}$. I.e. $\mathbf{W}$ is fixed and event $\mathcal{E}_\mathbf{W}$ in (6.2) is assumed to hold.

To be able to use the result of above proposition, we show that the minimizer of $\mathcal{L}(\boldsymbol{\theta})$ falls in $\mathcal{C}_\theta$ defined in (6.3).

Proposition 6.2 Assume $(\mathbf{x}, y)$ is generated according to (2.1), and recall that the rows of $\mathbf{W} \in \mathbb{R}^{N \times d}$ are drawn i.i.d. from $\text{Unif}(\mathbb{S}^{d-1})$. Let $\hat{\boldsymbol{\theta}} = \arg \min \mathcal{L}(\boldsymbol{\theta})$. We have $\hat{\boldsymbol{\theta}} \in \mathcal{C}_\theta$, with probability at least $1 - 4e^{-cn}$ for some absolute constant $c > 0$.

Motivated by Proposition 6.1 we define loss $\mathring{\mathcal{L}}(\boldsymbol{\theta})$ as follows:

$$(6.5) \quad \mathring{\mathcal{L}}(\boldsymbol{\theta}) := \frac{1}{2n} \sum_{i=1}^n \left(|y_i - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}_i)| + \varepsilon \|\mathbf{W}^\top \text{diag}(\mathbb{I}(\mathbf{W}\mathbf{x}_i > 0)) \boldsymbol{\theta}\|_{\ell_2} \right)^2 + \frac{\zeta}{2} \boldsymbol{\theta}^\top \boldsymbol{\Omega} \boldsymbol{\theta}.$$

By using Proposition 6.1, we can prove the following.

Proposition 6.3 Under the setting of Proposition 6.1 we have

$$(6.6) \quad \sup_{\boldsymbol{\theta} \in \mathcal{C}_\theta} \left| \mathcal{L}(\boldsymbol{\theta}) - \mathring{\mathcal{L}}(\boldsymbol{\theta}) \right| = o_{d, \mathbb{P}}(1).$$

6.2. Concentration of the adversarial effects As can be observed from the loss $\overset{\circ}{\mathcal{L}}(\boldsymbol{\theta})$, the effect of adversarial perturbation is reflected via the terms $\eta_i := \|\mathbf{W}^\top \text{diag}(\mathbb{I}(\mathbf{W}\mathbf{x}_i > 0)) \boldsymbol{\theta}\|_{\ell_2}$. In the next proposition, we show that apart from a negligible fraction of data points $i \in [n]$, the perturbation terms $\eta_i(\boldsymbol{\theta})^2$ concentrate around their expectation. All the proofs of the statements in this section are relegated to Appendix C.

Proposition 6.4 Let $\eta_i(\boldsymbol{\theta}) := \|\mathbf{W}^\top \text{diag}(\mathbb{I}(\mathbf{W}\mathbf{x}_i > 0)) \boldsymbol{\theta}\|_{\ell_2}$ and $\nu_i(\boldsymbol{\theta}, \gamma) := \mathbb{I}(|\eta_i(\boldsymbol{\theta})^2 - \mathbb{E}[\eta_i(\boldsymbol{\theta})^2]| > \gamma)$.

Under the setting of Proposition 6.1 and for any sequence γ_d such that $1/\gamma_d = e^{o(\sqrt{\log(d)})}$, we have

$$(6.7) \quad \sup_{\boldsymbol{\theta} \in \mathcal{C}_\theta} \frac{1}{n} \sum_{i=1}^n \nu_i(\boldsymbol{\theta}; \gamma_d) = O_{d, \mathbb{P}}(\log(d\gamma_d^2)^{-0.5}).$$

Corollary 6.5 By choosing sequence $\gamma_d = 1/\log(d)$ we obtain

$$(6.8) \quad \sup_{\boldsymbol{\theta} \in \mathcal{C}_\theta} \frac{1}{n} \sum_{i=1}^n \nu_i(\boldsymbol{\theta}; \frac{1}{\log(d)}) = o_{d, \mathbb{P}}(1).$$

By Corollary 6.5, other than at most an $o_d(1)$ fraction of data points $i \in [n]$, the terms $\eta_i(\boldsymbol{\theta})^2$ concentrate, in the sense $|\eta_i(\boldsymbol{\theta})^2 - \mathbb{E}[\eta_i(\boldsymbol{\theta})^2]| \leq 1/\log(d)$, uniformly over $\boldsymbol{\theta} \in \mathcal{C}_\theta$. This suggests that in the loss function we can replace the terms $\eta_i(\boldsymbol{\theta})$ by $\sqrt{\mathbb{E}[\eta_i(\boldsymbol{\theta})^2]}$. This observation will be formally stated in the next lemma. Before proceeding to it, let us compute the expectation of terms $\eta_i(\boldsymbol{\theta})^2$. We write

$$(6.9) \quad \mathbb{E}[\eta_i(\boldsymbol{\theta})^2] = \mathbb{E}\left[\|\mathbf{W}^\top \text{diag}(\mathbb{I}(\mathbf{W}\mathbf{x}_i > 0)) \boldsymbol{\theta}\|_{\ell_2}^2\right] = \boldsymbol{\theta}^\top \mathbb{E}\left[(\mathbf{W}\mathbf{W}^\top) \odot (\mathbb{I}(\mathbf{W}\mathbf{x}_i > 0)\mathbb{I}(\mathbf{W}\mathbf{x}_i > 0)^\top)\right] \boldsymbol{\theta},$$

where the expectation is with respect to $\mathbf{x}_i$ and $\mathbf{W}$ is fixed. Hence, this can be written as $\mathbb{E}[\eta_i(\boldsymbol{\theta})^2] = \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2$ with

$$\mathbf{J} := ((\mathbf{W}\mathbf{W}^\top) \odot \mathbb{E}[\mathbb{I}(\mathbf{W}\mathbf{x}_i > 0)\mathbb{I}(\mathbf{W}\mathbf{x}_i > 0)^\top])^{1/2}.$$

Note that the $\mathbf{J}$ is well-defined since the matrix under the square root is positive semidefinite. (This follows from the observation that the expression (6.9) is positive for all $\boldsymbol{\theta}$.)

Since $\|\mathbf{w}_\ell\|_{\ell_2} = 1$ and $\mathbf{x}_i \sim \mathcal{N}(0, \mathbf{I}_d)$, we have that $\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle$ and $\langle \mathbf{w}_k, \mathbf{x}_i \rangle$ are jointly Gaussian with

$$\mathbb{E}(\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle^2) = \mathbb{E}(\langle \mathbf{w}_k, \mathbf{x}_i \rangle^2) = 1, \quad \mathbb{E}(\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle \langle \mathbf{w}_k, \mathbf{x}_i \rangle) = \langle \mathbf{w}_k, \mathbf{w}_\ell \rangle,$$

and by using [15, Table 1] we get

$$\mathbb{E}[\mathbb{I}(\mathbf{W}\mathbf{x}_i > 0)\mathbb{I}(\mathbf{W}\mathbf{x}_i > 0)^\top] = \frac{\pi - \cos^{-1}(\mathbf{W}\mathbf{W}^\top)}{2\pi}.$$

Therefore, we obtain the following explicit formulation for $\mathbf{J}_\mathbf{W}$:

$$(6.10) \quad \mathbf{J} = \left((\mathbf{W}\mathbf{W}^\top) \odot \left(\frac{\pi - \cos^{-1}(\mathbf{W}\mathbf{W}^\top)}{2\pi} \right) \right)^{1/2}.$$

Motivated by Corollary 6.5 and the interpretation after it we define the loss function

$$(6.11) \quad \overset{\circ}{\mathcal{L}}(\boldsymbol{\theta}) := \frac{1}{2n} \sum_{i=1}^n (|y_i - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}_i)| + \varepsilon \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2})^2 + \frac{\zeta}{2} \boldsymbol{\theta}^\top \boldsymbol{\Omega} \boldsymbol{\theta}.$$

By our next lemma, the minimizer of $\overset{\circ}{\mathcal{L}}(\boldsymbol{\theta})$ converges to the minimizer of the original loss $\mathcal{L}(\boldsymbol{\theta})$ and therefore we can work with $\overset{\circ}{\mathcal{L}}(\boldsymbol{\theta})$ for our asymptotic analysis.

Lemma 6.6 *We have*

$$\sup_{\theta \in \mathcal{C}_\theta} \frac{|\overset{\circ\circ}{\mathcal{L}}(\theta) - \mathcal{L}(\theta)|}{1 + \min(\mathcal{L}(\theta), \overset{\circ\circ}{\mathcal{L}}(\theta))} = o_{d,\mathbb{P}}(1).$$

Also, by denoting by $\widehat{\theta}^*$ and $\widehat{\theta}$ the minimizers of $\overset{\circ\circ}{\mathcal{L}}(\theta)$ and $\mathcal{L}(\theta)$, we have $\|\widehat{\theta}^* - \widehat{\theta}\|_{\ell_2} \rightarrow 0$, in probability.

Motivated by the result of Lemma 6.6, we define a notion of adversarial risk based on the modified loss $\overset{\circ\circ}{\mathcal{L}}(\theta)$. Specifically, we define

$$(6.12) \quad \overset{\circ\circ}{\text{AR}}(\theta) := \mathbb{E}_{\mathbf{x}, y} \left[\left(|y - \theta^\top \sigma(\mathbf{W}\mathbf{x})| + \varepsilon_{\text{test}} \|\mathbf{J}\theta\|_{\ell_2} \right)^2 \right].$$

In the next lemma, we show that $\overset{\circ\circ}{\text{AR}}(\theta)$ converges to $\text{AR}(\theta)$ uniformly over $\mathcal{C}_\theta$.

Lemma 6.7 *Recall the adversarial risk of a model θ , denoted by $\text{AR}(\theta)$ and given by (4.2). Let $\overset{\circ\circ}{\text{AR}}(\theta)$ be defined as (6.12). We then have,*

$$\sup_{\theta \in \mathcal{C}_\theta} \frac{|\overset{\circ\circ}{\text{AR}}(\theta) - \text{AR}(\theta)|}{\sqrt{\text{AR}(\theta)}} = o_{d,\mathbb{P}}(1).$$

6.3. The Gaussian equivalence property and the noisy linear model In this section, we will show that in order to characterize the robust generalization error of the random features model, we can equivalently consider the so-called Gaussian features model (a.k.a. the noisy linear model). This equivalency is often termed as the *Gaussian Equivalence Property (GEP)*, and has recently been proven in several contexts [63, 39, 28, 19]. We prove this equivalency for adversarially-trained random features models in this section.

We begin with decomposing the nonlinear activation function $\sigma(z)$ as

$$(6.13) \quad \sigma(z) = \mu_0 + \mu_1 z + \mu_2 \sigma_\perp(z),$$

where for $G \sim \mathcal{N}(0, 1)$,

$$\mu_0 := \mathbb{E}[\sigma(G)], \quad \mu_1 = \mathbb{E}[G\sigma(G)], \quad \mu_2 := \sqrt{\mathbb{E}[\sigma^2(G)] - \mu_0^2 - \mu_1^2}.$$

For the case of shifted ReLU activation, defined in (4.1), we have $\mu_0 = 0$, $\mu_1 = \frac{1}{2}$ and $\mu_2 = \sqrt{\frac{1}{4} - \frac{1}{2\pi}}$. Also, $\sigma_\perp(z)$ is the nonlinear component of the activation function which is orthogonal to the constant and linear components in the following sense: $\mathbb{E}[\sigma_\perp(G)] = 0$ and $\mathbb{E}[G\sigma_\perp(G)] = 0$. We can then write the random features $\sigma(\mathbf{W}\mathbf{x})$ as follows

$$(6.14) \quad \sigma(\mathbf{W}\mathbf{x}) = \mu_0 \mathbf{1} + \mu_1 \mathbf{W}\mathbf{x} + \mu_2 \sigma_\perp(\mathbf{W}\mathbf{x}),$$

Note that the random variables $\sigma(\mathbf{w}_i^\top \mathbf{x})$ have zero mean and unit variance, by construction. Further, $\mathbb{E}_{\mathbf{x}}\{(\mathbf{w}_i^\top \mathbf{x})\sigma_\perp(\mathbf{w}_i^\top \mathbf{x})\} = 0$ since by construction $\mathbb{E}[\sigma_\perp(G)G] = 0$. This suggests to replace the variables $\sigma_\perp(\mathbf{w}_i^\top \mathbf{x})$ by a set of i.i.d standard normal variables and consider the following *noisy linear model*

$$(6.15) \quad \mathbf{f} := \mu_0 \mathbf{1} + \mu_1 \mathbf{W}\mathbf{x} + \mu_2 \mathbf{u},$$

with $\mathbf{f}, \mathbf{u} \in \mathbb{R}^N$ and $\mathbf{u} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_N)$ is generated independently from $\mathbf{x}$.

Consequently, we define the loss of the noisy linear model as

$$(6.16) \quad \mathcal{L}_{\text{nl}}(\boldsymbol{\theta}) := \frac{1}{2n} \sum_{i=1}^n (|y_i - \boldsymbol{\theta}^\top \mathbf{f}_i| + \varepsilon \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2})^2 + \frac{\zeta}{2} \boldsymbol{\theta}^\top \boldsymbol{\Omega} \boldsymbol{\theta},$$

where f_i are generated i.i.d. according to (6.15). We note that compared to the loss $\overset{\circ\circ}{\mathcal{L}}(\boldsymbol{\theta})$ defined in (6.11), we have only replaced the feature vectors $\sigma(\mathbf{W}\mathbf{x}_i)$ with the noisy linear features $\mathbf{f}_i$.

Let $\widehat{\boldsymbol{\theta}}^*$ and $\widehat{\boldsymbol{\theta}}_{\text{nl}}^*$ respectively denote the minimizers of $\overset{\circ\circ}{\mathcal{L}}(\boldsymbol{\theta})$ and $\mathcal{L}_{\text{nl}}(\boldsymbol{\theta})$. Roughly speaking, the Gaussian equivalence property (GEP) states that under certain conditions on $\mathbf{W}$ and the activation function σ , we have

$$(6.17) \quad \overset{\circ\circ}{\text{AR}}(\widehat{\boldsymbol{\theta}}^*) \approx \overset{\circ\circ}{\text{AR}}_{\text{nl}}(\widehat{\boldsymbol{\theta}}_{\text{nl}}^*) \quad \text{as } d \rightarrow \infty,$$

where $\overset{\circ\circ}{\text{AR}}(\cdot)$ is defined by (6.12) and $\overset{\circ\circ}{\text{AR}}_{\text{nl}}(\cdot)$ is its counterpart defined based on the noisy linear model, as follows:

$$(6.18) \quad \overset{\circ\circ}{\text{AR}}_{\text{nl}}(\boldsymbol{\theta}) := \mathbb{E}_{\mathbf{f}, y} \left[(|y - \boldsymbol{\theta}^\top \mathbf{f}| + \varepsilon_{\text{test}} \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2})^2 \right].$$

Therefore, by virtue of Lemma 6.7 and (6.17), we can henceforth focus on characterizing $\overset{\circ\circ}{\text{AR}}_{\text{nl}}(\widehat{\boldsymbol{\theta}}_{\text{nl}}^*)$.

In order to prove (6.17), we first show the asymptotic equality of $\overset{\circ\circ}{\text{AR}}_{\text{nl}}(\boldsymbol{\theta})$ and $\overset{\circ\circ}{\text{AR}}(\boldsymbol{\theta})$. All the proofs of the statements in this section are provided in Appendix D.

Proposition 6.8 *Consider model (2.1) under the asymptotic setting in Assumption 1 and define the set*

$$\mathcal{C}'_{\boldsymbol{\theta}} := \left\{ \boldsymbol{\theta} : \|\boldsymbol{\theta}\|_{\ell_\infty} \leq C_0 \sqrt{(\log(d))/d} \text{ and } \|\boldsymbol{\theta}\|_{\ell_2} \leq C_0 \text{ and } |\mathbf{1}^\top \boldsymbol{\theta}| \leq C_0 \sqrt{d/(\log(d))} \right\},$$

for an arbitrary but fixed constant $C_0 > 0$. Let $\overset{\circ\circ}{\text{AR}}(\boldsymbol{\theta})$ and $\overset{\circ\circ}{\text{AR}}_{\text{nl}}(\boldsymbol{\theta})$ be defined by (6.12)-(6.18). Then, for any $\boldsymbol{\theta} \in \mathcal{C}'_{\boldsymbol{\theta}}$ we have

$$(6.19) \quad \overset{\circ\circ}{\text{AR}}(\boldsymbol{\theta}) = \overset{\circ\circ}{\text{AR}}_{\text{nl}}(\boldsymbol{\theta}) + o_d(1),$$

In addition, we have the following characterizations for $\overset{\circ\circ}{\text{AR}}_{\text{nl}}(\boldsymbol{\theta})$:

$$(6.20) \quad \overset{\circ\circ}{\text{AR}}_{\text{nl}}(\boldsymbol{\theta}) = M(\boldsymbol{\theta})^2 + \varepsilon_{\text{test}}^2 \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2 + 2\sqrt{\frac{2}{\pi}} \varepsilon_{\text{test}} M(\boldsymbol{\theta}) \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2},$$

with $M(\boldsymbol{\theta})$ given by

$$(6.21) \quad M(\boldsymbol{\theta})^2 = \tau^2 + \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \left(\frac{1}{4} - \frac{1}{2\pi} \right) \|\boldsymbol{\theta}\|_{\ell_2}^2.$$

Proof of Proposition 6.8, i.e. equation (6.19), follows from a Central limit theorem (CLT) for weakly correlated variables proved in [30]. Specifically, [30] shows that $(\boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}), \boldsymbol{\beta}^\top \mathbf{x})$ converges in distribution to $(\boldsymbol{\theta}^\top \mathbf{f}, \boldsymbol{\beta}^\top \mathbf{x})$, where $\boldsymbol{\beta}$ is a fixed vector with bounded norm. In [39], the authors provide an alternative proof of this CLT using Stein's method and the Lindeberg approach. Their analysis assumes that the activation function $\sigma(z)$ is an odd function with bounded first derivatives. In addition, their analysis gives the convergence rate in terms of $\|\boldsymbol{\theta}\|_{\ell_\infty}$ (a Berry-Esseen type result).

By the characterizations (6.20), we know that, provided that $\theta \in \mathcal{C}'_\theta$, the quantity $\overset{\circ}{\text{AR}}_{\text{nl}}(\theta)$ depends on θ through the quantities $M(\theta)$ and $\|\mathbf{J}\theta\|_{\ell_2}$. As we show in Lemma F.3, $\|\mathbf{J}^2 - \mathbf{K}\| \rightarrow 0$, in probability with $\mathbf{K} = (\mathbf{W}\mathbf{W}^\top + \mathbf{I})/4$. Since by definition, for $\theta \in \mathcal{C}'_\theta$ we have $\|\theta\|_{\ell_2} \leq C_0$, therefore $\|\mathbf{J}\theta\|_{\ell_2}^2 \rightarrow (\|\mathbf{W}^\top \theta\|_{\ell_2}^2 + \|\theta\|_{\ell_2}^2)/4$. So, in order to show the GEP relation of the form (6.17), it suffices to show that the quantities $\|\theta\|_{\ell_2}$, $\|\frac{1}{2}\mathbf{W}^\top \theta - \beta\|_{\ell_2}$ and $\|\mathbf{W}^\top \theta\|_{\ell_2}$, evaluated at $\hat{\theta}^*$ converge to the corresponding quantities evaluated at $\hat{\theta}_{\text{nl}}^*$, and also $\hat{\theta}^*, \hat{\theta}_{\text{nl}}^* \in \mathcal{C}'_\theta$.

Theorem 6.9 *Consider the quantities Φ_A and Φ_B defined as*

$$\Phi_A := \min_{\theta} \frac{1}{n} \sum_{i=1}^n (|y_i - \theta^\top \sigma(\mathbf{W} \mathbf{x}_i)| + \varepsilon \|\mathbf{J}\theta\|_{\ell_2})^2 + \lambda \|\theta\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \theta - \beta \right\|_{\ell_2}^2 + \lambda_s \frac{\log(d)}{d} (\mathbf{1}^\top \theta)^2,$$

$$\Phi_B := \min_{\theta} \frac{1}{n} \sum_{i=1}^n (y_i - \theta^\top \mathbf{f}_i + \varepsilon \|\mathbf{J}\theta\|_{\ell_2})^2 + \lambda \|\theta\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \theta - \beta \right\|_{\ell_2}^2 + \lambda_s \frac{\log(d)}{d} (\mathbf{1}^\top \theta)^2,$$

where $\lambda, \lambda_s, \lambda_w > 0$, and $(\mathbf{x}_i, y_i)$ is generated i.i.d. according to (2.1). We further assume that the event $\mathcal{E}_W$ holds. Then, we have

$$(6.22) \quad \Phi_A \xrightarrow{\mathcal{P}} c \quad \text{if and only if} \quad \Phi_B \xrightarrow{\mathcal{P}} c,$$

where $\xrightarrow{\mathcal{P}}$ denotes convergence in probability.

Using this theorem, we can then prove the following proposition.

Proposition 6.10 *Recall $\hat{\theta}^*$ and $\hat{\theta}_{\text{nl}}^*$ given by*

$$\hat{\theta}^* = \arg \min_{\theta \in \mathbb{R}^N} \overset{\circ}{\mathcal{L}}(\theta) = \arg \min_{\theta \in \mathbb{R}^N} \frac{1}{2n} \sum_{i=1}^n (|y_i - \theta^\top \sigma(\mathbf{W} \mathbf{x}_i)| + \varepsilon \|\mathbf{J}\theta\|_{\ell_2})^2 + \frac{\zeta}{2} \theta^\top \Omega \theta,$$

$$(6.23) \quad \hat{\theta}_{\text{nl}}^* = \arg \min_{\theta \in \mathbb{R}^N} \mathcal{L}_{\text{nl}}(\theta) = \arg \min_{\theta \in \mathbb{R}^N} \frac{1}{2n} \sum_{i=1}^n (y_i - \theta^\top \mathbf{f}_i + \varepsilon \|\mathbf{J}\theta\|_{\ell_2})^2 + \frac{\zeta}{2} \theta^\top \Omega \theta.$$

Then, under the asymptotic regime of Assumption 1 we have

$$\hat{\theta}^*, \hat{\theta}_{\text{nl}}^* \in \mathcal{C}'_\theta,$$

with probability $1 - o_d(1)$, and

$$(6.24) \quad M(\hat{\theta}^*) - M(\hat{\theta}_{\text{nl}}^*) \xrightarrow{\mathcal{P}} 0, \quad \|\mathbf{J}\hat{\theta}^*\|_{\ell_2} - \|\mathbf{J}\hat{\theta}_{\text{nl}}^*\|_{\ell_2} \xrightarrow{\mathcal{P}} 0,$$

where $\xrightarrow{\mathcal{P}}$ denotes convergence in probability.

Therefore, the GEP (6.17) follows from combining Propositions 6.8 and 6.10.

Finally, the result of Theorem 4.2 follows by computing $\overset{\circ}{\text{AR}}_{\text{nl}}(\hat{\theta}_{\text{nl}}^*)$ when we send $\zeta \rightarrow 0$ after $d \rightarrow \infty$. This characterization will be carried out in Step 4 of the proof which will be described in the next section.

In non-adversarial contexts and for the standard risk (a.k.a. the generalization error) GEP has been observed by several previous works (see, e.g. [37, 63, 1, 61, 39, 31, 28, 30] and also [56, 12, 66] in the context of random kernel matrices). In [61] the authors provide a precise characterization of the standard risk for the random features model (in non-adversarial setting) and observed that it corresponds to that of its noisy linear counterpart model. A similar GEP phenomena was conjectured for maximum-margin linear classifiers in binary classification [63]. Subsequently, GEP has been proved for more general settings by [31] and [39] for a teacher-student framework. In [31] the authors show GEP for learning with one-pass

stochastic gradient descent (SGD). The work [39] considers the empirical risk minimization (with all data), which results in complicated correlations between the estimator and the samples, and proves the GEP for these settings. However, we cannot directly apply the result of [39] since it assumes that the activation function is an odd function, thrice continuously differentiable with bounded first three derivatives, which are violated for the ReLU activation. Also, our adversarial loss function has additional terms that are beyond the setting considered in [39]. Nevertheless, our proof of Theorem 6.9 is based on the machinery developed in [39]. Here we use a central limit theorem for weakly correlated random variables proved by [30] to show GEP in the context of adversarial training.

6.4. Analysis of the Gaussian noisy linear model via convex Gaussian minimax framework

In our final step, we provide a sharp characterization of the adversarial risk $\overset{\circ}{\text{AR}}_{\text{nl}}(\hat{\theta}_{\text{nl}}^*)$ using the Convex Gaussian Minimax Theorem (CGMT), which is a powerful and tight extension of Gordon’s Gaussian process inequality [33] with the presence of convexity. The underlying idea of the CGMT framework dates back to [80, 81, 82] where the constrained LASSO was analyzed in the high signal-to-noise ratio regimes. The seminal work [88, 87] significantly extended these ideas and developed the CGMT framework to precisely characterize the mean-squared errors of regularized M-estimators in high-dimensional linear models.

At a more technical level, the CGMT provides a principled machinery to characterize the asymptotic behavior of certain mini-max optimization problems that are affine in a Gaussian matrix $\mathbf{X}$, namely problems of the form

$$(6.25) \quad \min_{\theta} \max_u \quad u^\top \mathbf{X} \theta + \phi(\theta, u)$$

where $\phi(\theta, u)$ is convex in θ and concave in u . The CGMT decouples the above objective into a much simpler Gaussian process with essentially the same limit, yet much easier to analyze:

$$(6.26) \quad \min_{\theta} \max_u \quad \|\theta\|_{\ell_2} g^\top u + \|u\|_{\ell_2} h^\top \theta + \phi(\theta, u),$$

where g and h are independent Gaussian vectors with i.i.d $N(0, 1)$ entries. We refer to [88, Theorem 3] for a precise statement on the relation between the optimization (6.25), often referred to as *Primary Optimization (PO)* and (6.26), called *Auxiliary Optimization (AO)*. The next step is to derive the point-wise limit of the AO objective in the large dimension limit and showing that it concentrates around a deterministic function with a small number of scalar variables (called *scalarization* step). By showing that this convergence is uniform over a neighborhood of solution and using convexity-concavity of the function, we obtain a precise characterization of the adversarial risk in terms of the solutions of the corresponding convex-concave (deterministic) optimization (4.5).

Note that although the CGMT is a general machinery, the derivation and the study of the AO problem is entirely problem-specific and is usually rather challenging, often requiring the development of non-trivial probabilistic analysis. In relation to *Approximate message passing (AMP)*, which is another powerful tool for deriving asymptotically exact characterization of high-dimensional estimators (see e.g. [25]), it is worth noting that both of these techniques provide a deterministic equation (called state evolution in the AMP parlance) which describes the large limit behavior of a random system.

The CGMT has been recently used in several contexts, e.g., to characterize the performance of high-dimensional regularized logistic regression [73], SLOPE estimator in sparse linear regression [38], boosting and min ℓ_1 norm classifier [52], multi-class classification [89], and phase retrieval [20]. More closely to our work, the CGMT has been used to study the effect of adversarial training in the context of linear regression [46] and linear classifiers [45, 85]. On a technical side, the CGMT analysis for our current problem is more involved and intricate

than the analysis carried out in [46] for linear regression due to: (i) features $\mathbf{f}_i$ in (6.23) being correlated; (ii) the presence of the matrix $\mathbf{J}$ in the loss which introduces more interactions among the model parameters.

Acknowledgments The authors thank Alexander Robey for interesting discussions and feedback on an early draft.

Funding The research of H. Hassani is supported by the NSF CAREER award, AFOSR YIP, the Intel Rising Star award, as well as the AI Institute for Learning-Enabled Optimization at Scale (TILOS). A. Javanmard is partially supported by the Sloan Research Fellowship in mathematics, an Adobe Data Science Faculty Research Award, the NSF CAREER Award DMS-1844481 and NSF Award DMS-2311024.

SUPPLEMENTARY MATERIAL

Supplement to: “Precise Statistical Analysis of Classification Accuracies for Adversarial Training”

Due to space constraints, proofs of theorems and some of the technical details are provided in the Supplementary Material [36].

REFERENCES

- [1] ABBASI, E., SALEHI, F. and HASSIBI, B. (2019). Universality in Learning from Linear Measurements. *Advances in Neural Information Processing Systems* **32** 12372–12382.
- [2] BAI, Z. and SILVERSTEIN, J. W. (2010). *Spectral analysis of large dimensional random matrices* **20**. Springer.
- [3] BARTLETT, P. L., MONTANARI, A. and RAKHLIN, A. (2021). Deep learning: a statistical viewpoint. *arXiv preprint arXiv:2103.09177*.
- [4] BELKIN, M., HSU, D., MA, S. and MANDAL, S. (2019). Reconciling modern machine-learning practice and the classical bias–variance trade-off. *Proceedings of the National Academy of Sciences* **116** 15849–15854.
- [5] BELKIN, M., MA, S. and MANDAL, S. (2018). To Understand Deep Learning We Need to Understand Kernel Learning. In *International Conference on Machine Learning* 541–549.
- [6] BIGGIO, B., CORONA, I., MAIORCA, D., NELSON, B., ŠRNDIĆ, N., LASKOV, P., GIACINTO, G. and ROLI, F. (2013). Evasion attacks against machine learning at test time. In *Joint European conference on machine learning and knowledge discovery in databases* 387–402. Springer.
- [7] BILLINGSLEY, P. (1995). *Probability and Measure*. *Wiley Series in Probability and Statistics*. Wiley.
- [8] BOYD, S. and VANDENBERGHE, L. (2009). *Convex optimization*. Cambridge university press.
- [9] BUBECK, S., LI, Y. and NAGARAJ, D. M. (2021). A law of robustness for two-layers neural networks. In *Conference on Learning Theory* 804–820. PMLR.
- [10] BUBECK, S. and SELLKE, M. (2021). A Universal Law of Robustness via Isoperimetry. *Advances in Neural Information Processing Systems* **34**.
- [11] CARMON, Y., RAGHUNATHAN, A., SCHMIDT, L., LIANG, P. and DUCHI, J. C. (2019). Unlabeled data improves adversarial robustness. *arXiv preprint arXiv:1905.13736*.
- [12] CHENG, X. and SINGER, A. (2013). The spectrum of random inner-product kernel matrices. *Random Matrices: Theory and Applications* **2** 1350010.
- [13] COHEN, J., ROSENFELD, E. and KOLTER, Z. (2019). Certified adversarial robustness via randomized smoothing. In *International Conference on Machine Learning* 1310–1320. PMLR.
- [14] DANIELY, A. (2017). SGD learns the conjugate kernel class of the network. In *Advances in Neural Information Processing Systems* 2422–2430.
- [15] DANIELY, A., FROSTIG, R. and SINGER, Y. (2016). Toward deeper understanding of neural networks: the power of initialization and a dual view on expressivity. In *Proceedings of the 30th International Conference on Neural Information Processing Systems* 2261–2269.
- [16] DENG, Z., KAMMOUN, A. and THRAMOULIDIS, C. (2019). A Model of Double Descent for High-dimensional Binary Linear Classification. *arXiv preprint arXiv:1911.05822*.

- [17] DENG, Z., ZHANG, L., GHORBANI, A. and ZOU, J. (2021). Improving adversarial robustness via unlabeled out-of-domain data. In *International Conference on Artificial Intelligence and Statistics* 2845–2853. PMLR.
- [18] DEBERSIN, J. and LECUÉ, G. (2023). On the robustness to adversarial corruption and to heavy-tailed data of the Stahel–Donoho median of means. *Information and Inference: A Journal of the IMA* **12** 814–850.
- [19] DHIFALLAH, O. and LU, Y. M. (2020). A precise performance analysis of learning with random features. *arXiv preprint arXiv:2008.11904*.
- [20] DHIFALLAH, O., THRAMPOULIDIS, C. and LU, Y. M. (2018). Phase retrieval via polytope optimization: Geometry, phase transitions, and new algorithms. *arXiv preprint arXiv:1805.09555*.
- [21] DOBRIBAN, E., HASSANI, H., HONG, D. and ROBEY, A. (2020). Provable tradeoffs in adversarially robust classification. *arXiv preprint arXiv:2006.05161*.
- [22] DOHMATOB, E. (2021). Fundamental tradeoffs between memorization and robustness in random features and neural tangent regimes. *arXiv preprint arXiv:2106.02630*.
- [23] DONHAUSER, K., TIFREA, A., AERNI, M., HECKEL, R. and YANG, F. (2021). Interpolation can hurt robust generalization even when there is no noise. *Advances in Neural Information Processing Systems* **34**.
- [24] DONOHO, D. L., MALEKI, A. and MONTANARI, A. (2009). Message-passing algorithms for compressed sensing. *Proceedings of the National Academy of Sciences* **106** 18914–18919.
- [25] DONOHO, D. L., MALEKI, A. and MONTANARI, A. (2009). Message-passing algorithms for compressed sensing. *Proceedings of the National Academy of Sciences* **106** 18914–18919.
- [26] EL KAROUI, N. (2010). The spectrum of kernel random matrices. *The Annals of Statistics* **38** 1–50.
- [27] EL KAROUI, N. (2018). On the impact of predictor geometry on the performance on high-dimensional ridge-regularized generalized robust regression estimators. *Probability Theory and Related Fields* **170** 95–175.
- [28] GERACE, F., LOUREIRO, B., KRZAKALA, F., MÉZARD, M. and ZDEBOROVÁ, L. (2020). Generalisation error in learning with random features and the hidden manifold model. In *International Conference on Machine Learning* 3452–3462. PMLR.
- [29] GILMER, J., METZ, L., FAGHRI, F., SCHOENHOLZ, S. S., RAGHU, M., WATTENBERG, M. and GOODFELLOW, I. (2018). Adversarial spheres. *arXiv preprint arXiv:1801.02774*.
- [30] GOLDT, S., LOUREIRO, B., REEVES, G., KRZAKALA, F., MÉZARD, M. and ZDEBOROVÁ, L. (2020). The Gaussian equivalence of generative models for learning with shallow neural networks. *arXiv preprint arXiv:2006.14709*.
- [31] GOLDT, S., MÉZARD, M., KRZAKALA, F. and ZDEBOROVÁ, L. (2020). Modeling the influence of data structure on learning in neural networks: The hidden manifold model. *Physical Review X* **10** 041044.
- [32] GOODFELLOW, I. J., SHLENS, J. and SZEGEDY, C. (2015). Explaining and Harnessing Adversarial Examples. In *3rd International Conference on Learning Representations, ICLR 2015, San Diego, CA, USA, May 7-9, 2015, Conference Track Proceedings*.
- [33] GORDON, Y. (1988). On Milman’s inequality and random subspaces which escape through a mesh in $\mathbb{R}^n$. In *Geometric aspects of functional analysis* 84–106. Springer.
- [34] GOWAL, S., QIN, C., UESATO, J., MANN, T. and KOHLI, P. (2020). Uncovering the Limits of Adversarial Training against Norm-Bounded Adversarial Examples. *arXiv preprint arXiv:2010.03593*.
- [35] GUNASEKAR, S., LEE, J. D., SOUDRY, D. and SREBRO, N. (2018). Implicit Bias of Gradient Descent on Linear Convolutional Networks. In *Advances in Neural Information Processing Systems* (S. BENGIO, H. WALLACH, H. LAROCHELLE, K. GRAUMAN, N. CESA-BIANCHI and R. GARNETT, eds.) **31** 9461–9471. Curran Associates, Inc.
- [36] HASSANI, H. and JAVANMARD, A. (2022). Supplementary material to “The curse of overparametrization in adversarial training: Precise analysis of robust generalization for random features regression”.
- [37] HASTIE, T., MONTANARI, A., ROSSET, S. and TIBSHIRANI, R. J. (2022). Surprises in high-dimensional ridgeless least squares interpolation. *The Annals of Statistics* **50** 949–986.
- [38] HU, H. and LU, Y. M. (2019). Asymptotics and optimal designs of SLOPE for sparse linear regression. In *2019 IEEE International Symposium on Information Theory (ISIT)* 375–379. IEEE.
- [39] HU, H. and LU, Y. M. (2020). Universality laws for high-dimensional learning with random features. *arXiv preprint arXiv:2009.07669*.
- [40] HUANG, S.-T. and LEDERER, J. (2023). DeepMoM: Robust Deep Learning With Median-of-Means. *Journal of Computational and Graphical Statistics* **32** 181–195.
- [41] JACOT, A., GABRIEL, F. and HONGLER, C. (2018). Neural Tangent Kernel: Convergence and Generalization in Neural Networks. In *Advances in Neural Information Processing Systems* (S. BENGIO, H. WALLACH, H. LAROCHELLE, K. GRAUMAN, N. CESA-BIANCHI and R. GARNETT, eds.) **31** 8571–8580. Curran Associates, Inc.
- [42] JACOT, A., GABRIEL, F. and HONGLER, C. (2018). Neural tangent kernel: Convergence and generalization in neural networks. In *Advances in neural information processing systems* 8571–8580.

- [43] JALAL, A., ILYAS, A., DASKALAKIS, C. and DIMAKIS, A. G. (2017). The robust manifold defense: Adversarial training using generative models. *arXiv preprint arXiv:1712.09196*.
- [44] JAVANMARD, A., MONDELLI, M. and MONTANARI, A. (2020). Analysis of a two-layer neural network via displacement convexity. *The Annals of Statistics* **48** 3619–3642.
- [45] JAVANMARD, A. and SOLTANOLKOTABI, M. (2022). Precise statistical analysis of classification accuracies for adversarial training. *The Annals of Statistics* **50** 2127–2156.
- [46] JAVANMARD, A., SOLTANOLKOTABI, M. and HASSANI, H. (2020). Precise tradeoffs in adversarial training for linear regression. In *Conference on Learning Theory* 2034–2078. PMLR.
- [47] KURAKIN, A., GOODFELLOW, I. and BENGIO, S. (2016). Adversarial machine learning at scale. *arXiv preprint arXiv:1611.01236*.
- [48] LAI, L. and BAYRAKTAR, E. (2020). On the adversarial robustness of robust estimators. *IEEE Transactions on Information Theory* **66** 5097–5109.
- [49] LAURENT, B. and MASSART, P. (2000). Adaptive estimation of a quadratic functional by model selection. *Annals of Statistics* 1302–1338.
- [50] LEDOUX, M. (2001). The concentration of measure phenomenon. *volume 89 of Mathematical Surveys and Monographs. American Mathematical Society, Providence, RI*.
- [51] LI, Y. and LIANG, Y. (2018). Learning overparameterized neural networks via stochastic gradient descent on structured data. *NeurIPS*.
- [52] LIANG, T. and SUR, P. (2020). A Precise High-Dimensional Asymptotic Theory for Boosting and Min-L1-Norm Interpolated Classifiers. *arXiv preprint arXiv:2002.01586*.
- [53] LIESE, F. and MIESCKE, K.-J. (2008). Statistical decision theory. In *Statistical Decision Theory: Estimation, Testing, and Selection* 1–52. Springer.
- [54] LIESE, F. and MIESCKE, K.-J. (2008). Statistical Decision Theory: Estimation, Testing, and Selection. In *Springer Science & Business Media*.
- [55] LINDBERG, J. W. (1922). Eine neue Herleitung des Exponentialgesetzes in der Wahrscheinlichkeitsrechnung. *Mathematische Zeitschrift* **15** 211–225.
- [56] LOUART, C., LIAO, Z. and COUILLET, R. (2018). A random matrix approach to neural networks. *The Annals of Applied Probability* **28** 1190–1248.
- [57] MADRY, A., MAKELOV, A., SCHMIDT, L., TSIPRAS, D. and VLADU, A. (2018). Towards Deep Learning Models Resistant to Adversarial Attacks. In *6th International Conference on Learning Representations, ICLR 2018, Vancouver, BC, Canada, April 30 - May 3, 2018, Conference Track Proceedings*.
- [58] MADRY, A., MAKELOV, A., SCHMIDT, L., TSIPRAS, D. and VLADU, A. (2018). Towards Deep Learning Models Resistant to Adversarial Attacks. In *International Conference on Learning Representations*.
- [59] MAHLOUJIFAR, S., DIOCHNOS, D. I. and MAHMOODY, M. (2019). The curse of concentration in robust learning: Evasion and poisoning attacks from concentration of measure. In *Proceedings of the AAAI Conference on Artificial Intelligence* **33** 4536–4543.
- [60] MAHLOUJIFAR, S. and MAHMOODY, M. (2019). Can Adversarially Robust Learning Leverage Computational Hardness? In *Algorithmic Learning Theory* 581–609. PMLR.
- [61] MEI, S. and MONTANARI, A. (2021). The generalization error of random features regression: Precise asymptotics and double descent curve. *Communications on Pure and Applied Mathematics*, doi.org/10.1002/cpa.22008.
- [62] MIN, Y., CHEN, L. and KARBASI, A. (2021). The curious case of adversarially robust models: More data can help, double descend, or hurt generalization. In *Uncertainty in Artificial Intelligence* 129–139. PMLR.
- [63] MONTANARI, A., RUAN, F., SOHN, Y. and YAN, J. (2019). The generalization error of max-margin linear classifiers: High-dimensional asymptotics in the overparametrized regime. *arXiv preprint arXiv:1911.01544*.
- [64] MONTANARI, A., ZHONG, Y. and ZHOU, K. (2021). Tractability from overparametrization: The example of the negative perceptron. *arXiv preprint arXiv:2110.15824*.
- [65] NAJAFI, A., MAEDA, S.-I., KOYAMA, M. and MIYATO, T. (2019). Robustness to adversarial perturbations in learning from incomplete data. *arXiv preprint arXiv:1905.13021*.
- [66] PENNINGTON, J. and WORAHL, P. (2019). Nonlinear random matrix theory for deep learning. *Journal of Statistical Mechanics: Theory and Experiment* **2019** 124005.
- [67] RAGHUNATHAN, A., XIE, S. M., YANG, F., DUCHI, J. C. and LIANG, P. (2019). Adversarial Training Can Hurt Generalization. *arXiv preprint arXiv:1906.06032*.
- [68] RAHIMI, A. and RECHT, B. (2007). Random features for large-scale kernel machines. *Advances in neural information processing systems* **20** 1177–1184.
- [69] RAHIMI, A. and RECHT, B. (2008). Uniform approximation of functions with random bases. In *2008 46th Annual Allerton Conference on Communication, Control, and Computing* 555–561. IEEE.
- [70] REBUFFI, S.-A., GOWAL, S., CALIAN, D. A., STIMBERG, F., WILES, O. and MANN, T. A. (2021). Data Augmentation Can Improve Robustness. *Advances in Neural Information Processing Systems* **34**.

- [71] RICHARDSON, T. and URBANKE, R. (2008). *Modern coding theory*. Cambridge university press.
- [72] RUDELSON, M., VERSHYNIN, R. et al. (2013). Hanson-Wright inequality and sub-gaussian concentration. *Electronic Communications in Probability* **18**.
- [73] SALEHI, F., ABBASI, E. and HASSIBI, B. (2019). The Impact of Regularization on High-dimensional Logistic Regression. In *Advances in Neural Information Processing Systems* (H. WALLACH, H. LAROCHELLE, A. BEYGEZIMER, F. D'ALCHÉ-BUC, E. FOX and R. GARNETT, eds.) **32**. Curran Associates, Inc.
- [74] SEHWAG, V., MAHLOUJIFAR, S., HANDINA, T., DAI, S., XIANG, C., CHIANG, M. and MITTAL, P. (2021). Improving Adversarial Robustness Using Proxy Distributions. *arXiv preprint arXiv:2104.09425*.
- [75] SHAFABI, A., HUANG, W. R., STUDER, C., FEIZI, S. and GOLDSTEIN, T. (2019). Are adversarial examples inevitable? In *7th International Conference on Learning Representations, ICLR 2019, New Orleans, LA, USA, May 6-9, 2019*.
- [76] SION, M. et al. (1958). On general minimax theorems. *Pacific Journal of mathematics* **8** 171–176.
- [77] SOLTANOLKOTABI, M., JAVANMARD, A. and LEE, J. D. (2018). Theoretical insights into the optimization landscape of over-parameterized shallow neural networks. *IEEE Transactions on Information Theory* **65** 742–769.
- [78] SONG, M., MONTANARI, A. and NGUYEN, P. (2018). A mean field view of the landscape of two-layers neural networks. In *Proceedings of the National Academy of Sciences* **115** E7665–E7671.
- [79] SOUDRY, D., HOFFER, E., NACSON, M. S., GUNASEKAR, S. and SREBRO, N. (2018). The implicit bias of gradient descent on separable data. *The Journal of Machine Learning Research* **19** 2822–2878.
- [80] STOJNIC, M. (2013). A framework to characterize performance of LASSO algorithms. *arXiv preprint arXiv:1303.7291*.
- [81] STOJNIC, M. (2013). Meshes that trap random subspaces. *arXiv preprint arXiv:1304.0003*.
- [82] STOJNIC, M. (2013). Upper-bounding ℓ_1 -optimization weak thresholds. *arXiv preprint arXiv:1303.7289*.
- [83] SU, D., ZHANG, H., CHEN, H., YI, J., CHEN, P.-Y. and GAO, Y. (2018). Is Robustness the Cost of Accuracy?—A Comprehensive Study on the Robustness of 18 Deep Image Classification Models. In *Proceedings of the European Conference on Computer Vision (ECCV)* 631–648.
- [84] SZEGEDY, C., ZAREMBA, W., SUTSKEVER, I., BRUNA, J., ERHAN, D., GOODFELLOW, I. J. and FERGUS, R. (2014). Intriguing properties of neural networks. ICLR, abs/1312.6199, 2014.
- [85] TAHERI, H., PEDARSANI, R. and THRAMOULIDIS, C. (2020). Asymptotic behavior of adversarial training in binary classification. *arXiv preprint arXiv:2010.13275*.
- [86] THRAMOULIDIS, C., ABBASI, E. and HASSIBI, B. (2015). Precise high-dimensional error analysis of regularized m-estimators. In *2015 53rd Annual Allerton Conference on Communication, Control, and Computing (Allerton)* 410–417. IEEE.
- [87] THRAMOULIDIS, C., ABBASI, E. and HASSIBI, B. (2018). Precise Error Analysis of Regularized M -Estimators in High Dimensions. *IEEE Transactions on Information Theory* **64** 5592–5628.
- [88] THRAMOULIDIS, C., OYMAK, S. and HASSIBI, B. (2015). Regularized linear regression: A precise analysis of the estimation error. In *Conference on Learning Theory* 1683–1709.
- [89] THRAMOULIDIS, C., OYMAK, S. and SOLTANOLKOTABI, M. (2020). Theoretical insights into multiclass classification: A high-dimensional asymptotic view. *arXiv preprint arXiv:2011.07729*.
- [90] TSIPRAS, D., SANTURKAR, S., ENGSTROM, L., TURNER, A. and MADRY, A. (2019). Robustness May Be at Odds with Accuracy. In *International Conference on Learning Representations*.
- [91] VERSHYNIN, R. (2010). Introduction to the non-asymptotic analysis of random matrices. *arXiv preprint arXiv:1011.3027*.
- [92] VERSHYNIN, R. (2018). *High-dimensional probability: An introduction with applications in data science* **47**. Cambridge University Press.
- [93] WONG, E. and KOLTER, J. Z. (2018). Provable Defenses against Adversarial Examples via the Convex Outer Adversarial Polytope. In *Proceedings of the 35th International Conference on Machine Learning, ICML 2018, Stockholmsmässan, Stockholm, Sweden, July 10-15, 2018* 5283–5292.
- [94] WU, B., CHEN, J., CAI, D., HE, X. and GU, Q. (2021). Do Wider Neural Networks Really Help Adversarial Robustness? *Advances in Neural Information Processing Systems* **34**.
- [95] ZHAI, R., CAI, T., HE, D., DAN, C., HE, K., HOPCROFT, J. and WANG, L. (2019). Adversarially robust generalization just requires more unlabeled data. *arXiv preprint arXiv:1906.00555*.
- [96] ZHANG, H., WU, Y. and HUANG, H. (2022). How Many Data Are Needed for Robust Learning? *arXiv preprint arXiv:2202.11592*.
- [97] ZHANG, H., YU, Y., JIAO, J., XING, E. P., GHAOUI, L. E. and JORDAN, M. I. (2019). Theoretically Principled Trade-off between Robustness and Accuracy. In *Proceedings of the 36th International Conference on Machine Learning, ICML 2019, 9-15 June 2019, Long Beach, California, USA* 7472–7482.

SUPPLEMENTARY MATERIAL TO “PRECISE STATISTICAL ANALYSIS OF CLASSIFICATION ACCURACIES FOR ADVERSARIAL TRAINING”

BY HAMED HASSANI^{1,a}, ADEL JAVANMARD^{2,b} 

¹*Department of Electrical and Systems Engineering, University of Pennsylvania, hassani@seas.upenn.edu*

²*Data Sciences and Operations Department, University of Southern California, ajavanma@usc.edu*

The supplementary materials contain the proofs of theorems and technical lemmas. It is structured around the main four steps outlined in Section 6.

For the sake of completeness, we reintroduce the notation used throughout the proofs.

Notations. Throughout the paper, we use $O_d(\cdot)$, $o_d(\cdot)$ to denote the standard big-O and little-o notation, where we stress the asymptotic variable d . Likewise, we denote by $O_{d,\mathbb{P}}$ and $o_{d,\mathbb{P}}$ to indicate asymptotic behavior in probability. Specifically, $f(d) = O_{d,\mathbb{P}}(g(d))$ if for any $\varepsilon > 0$, there exists $C_\varepsilon > 0$ and large enough d_ε such that $\mathbb{P}(|f(d)/g(d)| > C_\varepsilon) \leq \varepsilon$, for all $d \geq d_\varepsilon$. Similarly, $f(d) = o_{d,\mathbb{P}}(g(d))$ if $f(d)/g(d)$ converges to zero in probability. We write $f(d) \approx g(d)$ as $d \rightarrow \infty$, when $f(d) - g(d) \rightarrow 0$, in probability. Note that we consider the asymptotic regime where n, d, N grow at the same scale, ($\lim N/d \rightarrow \psi_1$ and $\lim n/d \rightarrow \psi_2$ for some positive constants ψ_1 and ψ_2), the expression $d \rightarrow \infty$ implies that $n, N \rightarrow \infty$, as well.

For a matrix $\mathbf{A}$, we denote by $\|\mathbf{A}\|$ its operator norm, $\|\mathbf{A}\|_F = (\sum_{ij} A_{ij}^2)^{1/2}$ the Frobenius norm of $\mathbf{A}$. For two matrices $\mathbf{A}$ and $\mathbf{B}$ of same size, we let $\mathbf{A} \odot \mathbf{B}$ be the element-wise product of $\mathbf{A}$ and $\mathbf{B}$. In addition, $[\mathbf{A}; \mathbf{B}]$ concatenates the two matrices row-wise and $[\mathbf{A}, \mathbf{B}]$ denotes the column-wise concatenation. For an integer n , we use the shorthand $[n] = \{1, \dots, n\}$.

CONTENTS

A	Interchanging the limits of $d \rightarrow \infty$ and $\zeta \rightarrow 0$	2
A.1	Proof of Lemma A.3	4
B	Proofs of step 1: Asymptotically-exact closed form of adversarial examples	4
B.1	Proof of Lemma B.1	6
B.2	Proof of Proposition 6.2	7
B.3	Proof of Proposition 6.3	11
C	Proofs of step 2: Concentration of the adversarial effects	12
C.1	Proof of Proposition 6.4	12
C.2	Proof of Lemma C.1	13
C.3	Proof of Lemma C.2	16
C.4	Proof of Lemma 6.6	17
C.5	Proof of Lemma 6.7	18
D	Proofs of step 3: The Gaussian equivalence property	19
D.1	Proof of Proposition 6.8	19
D.2	Proof of Theorem 6.9	21
D.3	Proof of Proposition 6.10	44
E	Proofs of Step 4: Analysis of the Gaussian noisy linear model via convex Gaussian minimax framework	45
E.1	Scalarization of the AO problem	47
E.2	Convergence analysis of the AO problem	51

E.3	Proof of Theorem 4.2(b)	55
E.4	Proof of Proposition 5.1	57
E.5	Proofs of the Auxiliary Lemmas	58
F	Some useful lemmas	60

APPENDIX A: INTERCHANGING THE LIMITS OF $d \rightarrow \infty$ AND $\zeta \rightarrow 0$

Consider the loss function (5.1) given by

$$\mathcal{L}(\boldsymbol{\theta}, \zeta, d) = \max_{\|\boldsymbol{\delta}_i\|_{\ell_2} \leq \epsilon} \frac{1}{2n} \sum_{i=1}^n \left(y_i - \boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i)) \right)^2 + \frac{\zeta}{2} \boldsymbol{\theta}^\top \boldsymbol{\Omega} \boldsymbol{\theta},$$

where with a slight abuse of notation, we made the dependence on ζ and d explicit.

In the next lemma we show that the order of the two limits $d \rightarrow \infty$ and $\zeta \rightarrow 0$ can be interchanged.

Lemma A.1 *Under the assumptions of Theorem 4.2, we have*

$$\lim_{d \rightarrow \infty} \min_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}, 0, d) = \lim_{\zeta \rightarrow 0} \lim_{d \rightarrow \infty} \min_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}, \zeta, d).$$

Proof (Proof of Lemma A.1) First note that

$$(A.1) \quad \min_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}, 0, d) = \min_{\boldsymbol{\theta}, \zeta \geq 0} \mathcal{L}(\boldsymbol{\theta}, \zeta, d) = \min_{\zeta \geq 0} \min_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}, \zeta, d) = \lim_{\zeta \rightarrow 0} \min_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}, \zeta, d).$$

The last step holds since $\mathcal{L}(\boldsymbol{\theta}, \zeta, d)$ is increasing in ζ for all $\boldsymbol{\theta}$:

$$\mathcal{L}(\boldsymbol{\theta}, \zeta_1, d) \leq \mathcal{L}(\boldsymbol{\theta}, \zeta_2, d), \quad \text{if } \zeta_1 \leq \zeta_2.$$

Minimizing both sides over $\boldsymbol{\theta}$, we get that $\min_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}, \zeta, d)$ is increasing in ζ .

We next show that

$$(A.2) \quad \lim_{d \rightarrow \infty} \lim_{\zeta \rightarrow 0} \min_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}, \zeta, d) = \lim_{\zeta \rightarrow 0} \lim_{d \rightarrow \infty} \min_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}, \zeta, d),$$

where the limits are in probability. Without loss of generality we restrict the domain of ζ to $[0, \zeta_*]$, for an arbitrary but fixed ζ_* . The reason is that in our proofs provided in the paper we allow ζ to be an arbitrarily small fixed value (i.e. we need ζ to be arbitrarily small, but fixed). We next use the Moore-Osgood theorem on exchanging limits, by which we need to verify that

$$(A.3) \quad \lim_{d \rightarrow \infty} \min_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}, \zeta, d) = f(\zeta), \quad \text{uniformly on } \zeta \in (0, \zeta_*],$$

$$(A.4) \quad \lim_{\zeta \rightarrow 0} \min_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}, \zeta, d) = A_d, \quad \text{pointwise over } d \in \mathbb{N}.$$

The second identity follows from (A.1). To prove the first identity, note that $\mathcal{L}(\boldsymbol{\theta}, \zeta, d)$ is convex in $(\boldsymbol{\theta}, \zeta)$. Now, since partial minimization preserves convexity [8, Section 3.2.5], $\min_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}, \zeta, d)$ is convex in ζ . The point-wise limit of (A.3) is already established in the paper, and we obtain uniform convergence using the convexity lemma [53, Lemma 7.75]. In words, the lemma states that pointwise convergence of convex functions implies uniform convergence in compact subsets.

Combining (A.1) and (A.2) we obtain

$$(A.5) \quad \lim_{d \rightarrow \infty} \min_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}, 0, d) = \lim_{\zeta \rightarrow 0} \lim_{d \rightarrow \infty} \min_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}, \zeta, d).$$

■

Recall that our main goal in the paper is to characterize the in-probability limit of $\text{AR}(\widehat{\theta}^\varepsilon)$, with

$$(A.6) \quad \widehat{\theta}^\varepsilon = \arg \min_{\theta} \mathcal{L}(\theta, 0, d).$$

Define

$$(A.7) \quad \widehat{\theta}_\zeta^\varepsilon = \arg \min_{\theta} \mathcal{L}(\theta, \zeta, d).$$

Our next lemma relates $\text{AR}(\widehat{\theta}^\varepsilon)$ to $\text{AR}(\widehat{\theta}_\zeta^\varepsilon)$.

Proposition A.2 *Let $\widehat{\theta}^\varepsilon$ and $\widehat{\theta}_\zeta^\varepsilon$ be respectively given by (A.6) and (A.7). Under assumptions of Theorem 4.2, we have*

$$\lim_{\zeta \rightarrow 0} \lim_{d \rightarrow \infty} \text{AR}(\widehat{\theta}_\zeta^\varepsilon) = \lim_{d \rightarrow \infty} \text{AR}(\widehat{\theta}^\varepsilon).$$

Proof (Proof of Proposition A.2) To proof the claim, we use a standard trick to translate the question on the optimal solution of the minimization problem (i.e. $\widehat{\theta}^\varepsilon, \widehat{\theta}_\zeta^\varepsilon$) to one regarding the optimal costs.

Let $B = \lim_{\zeta \rightarrow 0} \lim_{d \rightarrow \infty} \text{AR}(\widehat{\theta}_\zeta^\varepsilon)$ (which exists and is calculated in Theorem 4.2). We need to show that $\widehat{\theta}^\varepsilon$ belongs to the following set $S_\delta := \{\theta : |\text{AR}(\theta) - B| \leq \delta\}$, with probability converging to one (as $d \rightarrow \infty$) for all $\delta > 0$. Let S_δ^c denotes the complement set. If we show that

$$(A.8) \quad \min_{\theta \in S_\delta^c} \mathcal{L}(\theta, 0, d) > \mathcal{L}(\widehat{\theta}^\varepsilon, 0, d),$$

then $\widehat{\theta}^\varepsilon$ must lie in S_δ . We formalize it in the next lemma.

Lemma A.3 *Suppose that there exist constants $\ell, \tilde{\ell}$ and $\eta > 0$ such that*

- $\tilde{\ell} \geq \ell + 2\eta$,
- $\mathcal{L}(\widehat{\theta}^\varepsilon, 0, d) < \ell + \eta$ with probability at least $1 - p$,
- $\min_{\theta \in S_\delta^c} \mathcal{L}(\theta, 0, d) > \tilde{\ell} - \eta$ with probability at least $1 - p$.

Then, $\mathbb{P}(\widehat{\theta}^\varepsilon \in S_\delta) \geq 1 - 2p$.

We then have the following corollary.

Corollary A.4 *Suppose that there exist constants $\ell < \tilde{\ell}$ such that $\mathcal{L}(\widehat{\theta}^\varepsilon, 0, d) \xrightarrow{p} \ell$ and $\min_{\theta \in S_\delta^c} \mathcal{L}(\theta, 0, d) \xrightarrow{p} \tilde{\ell}$. Then, $\lim_{d \rightarrow \infty} \mathbb{P}(\widehat{\theta}^\varepsilon \in S_\delta) = 1$, for every $\delta > 0$.*

In light of the above corollary, we compare the converging limits: Let $\ell := \lim_{d \rightarrow \infty} \mathcal{L}(\widehat{\theta}^\varepsilon, 0, d)$ and $\tilde{\ell} := \lim_{d \rightarrow \infty} \min_{\theta \in S_\delta^c} \mathcal{L}(\theta, 0, d)$. We need to show $\ell < \tilde{\ell}$. A similar trick has been used in [86, Theorem 6.1 (iii)]. We next use (A.5), by which

$$\ell = \lim_{\zeta \rightarrow 0} \lim_{d \rightarrow \infty} \min_{\theta} \mathcal{L}(\theta, \zeta, d).$$

By a similar argument,

$$\tilde{\ell} = \lim_{\zeta \rightarrow 0} \lim_{d \rightarrow \infty} \min_{\theta \in S_\delta^c} \mathcal{L}(\theta, \zeta, d),$$

where the difference is the domain over which we optimize. Note that in Theorem 4.2 we calculate ℓ as the optimal value of a deterministic convex-concave optimization problem and show that it has a unique solution. Likewise, one can obtain a similar optimization for $\tilde{\ell}$ with the difference that its variables come from a restricted domain, which excludes the optimal solution of the former. By uniqueness of the solution, we conclude that $\ell < \tilde{\ell}$, which completes the proof of proposition. $\blacksquare$

A.1. Proof of Lemma A.3

Define the event

$$\mathcal{E} := \left\{ \min_{\boldsymbol{\theta} \in S_\delta^c} \mathcal{L}(\boldsymbol{\theta}, 0, d) > \tilde{\ell} - \eta, \mathcal{L}(\widehat{\boldsymbol{\theta}}^\varepsilon, 0, d) < \ell + \eta \right\}.$$

On this event, using the first condition, we see that (A.8) holds and so $\widehat{\boldsymbol{\theta}}^\varepsilon \in S_\delta$. So we need to show that $\mathbb{P}(\mathcal{E}) \geq 1 - 2p$, which follows easily from union bounding and using the second and third conditions.

Remark A.2 In [61] the authors derive a precise characterization of the generalization of random features regression in a non-adversarial setting. This work makes a conjecture (see Remark 1 therein) that the generalization error of ridgeless estimator is the same as the min-norm least square estimator. This conjecture amounts to showing that the limits $\lambda \rightarrow 0$ (λ the ridge penalty parameter) and $d \rightarrow \infty$ can be exchanged. We believe that this conjecture can be proved by following a similar argument of the proof of Lemma A.1.

APPENDIX B: PROOFS OF STEP 1: ASYMPTOTICALLY-EXACT CLOSED FORM OF ADVERSARIAL EXAMPLES

Recall that $\mathbf{x} \sim \mathcal{N}(0, \mathbf{I}_d)$. In the following we will show that conditioned on the event $\mathcal{E}_{\mathbf{W}}$, defined in (6.2), with probability at least $1 - c/(\log(d))^2 - N^2 d^{-C}$ over the choice of $\mathbf{x}$, we have

$$(B.1) \quad \sup_{\boldsymbol{\theta} \in \mathcal{C}_\theta, \|\boldsymbol{\delta}\|_{\ell_2} \leq \varepsilon} \left| \boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x} + \boldsymbol{\delta})) - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}) - \boldsymbol{\theta}^\top \text{diag}(\mathbb{I}(\mathbf{W}\mathbf{x} > 0)) \mathbf{W} \boldsymbol{\delta} \right| = C \frac{\log(d)}{d^{\frac{1}{6}}}.$$

As a result, we can write

$$(B.2) \quad \max_{\|\boldsymbol{\delta}\|_{\ell_2} \leq \varepsilon} |y - \boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x} + \boldsymbol{\delta}))| = \max_{\|\boldsymbol{\delta}\|_{\ell_2} \leq \varepsilon} |y - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}) - \langle \mathbf{W}^\top \text{diag}(\mathbb{I}(\mathbf{W}\mathbf{x} > 0)) \boldsymbol{\theta}, \boldsymbol{\delta} \rangle| + C \frac{\log(d)}{d^{\frac{1}{6}}},$$

for an absolute constant $C > 0$, uniformly over $\boldsymbol{\theta} \in \mathcal{C}_\theta$. The maximization problem in the right-hand side of the above relation has a closed-form solution:

$$(B.3) \quad \boldsymbol{\delta} = \varepsilon \text{sign}(y - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x})) \frac{\mathbf{W}^\top \text{diag}(\mathbb{I}(\mathbf{W}\mathbf{x} > 0)) \boldsymbol{\theta}}{\|\mathbf{W}^\top \text{diag}(\mathbb{I}(\mathbf{W}\mathbf{x} > 0)) \boldsymbol{\theta}\|_{\ell_2}},$$

which gives us the desired result (6.4). It thus remains to prove (B.1).

Denote the rows of matrix $\mathbf{W}$ by $\{\mathbf{w}_1, \dots, \mathbf{w}_N\}$ with $\mathbf{w}_\ell \in \mathbb{R}^d$. Given $\mathbf{x}$, we define the three sets

$$A(\mathbf{x}) = \left\{ \ell : \langle \mathbf{w}_\ell, \mathbf{x} \rangle > d^{-\frac{1}{3}} \right\},$$

$$B(\mathbf{x}) = \left\{ \ell : \langle \mathbf{w}_\ell, \mathbf{x} \rangle < -d^{-\frac{1}{3}} \right\},$$

$$C(\mathbf{x}) = \left\{ \ell : \langle \mathbf{w}_\ell, \mathbf{x} \rangle \in [-d^{-\frac{1}{3}}, d^{-\frac{1}{3}}] \right\}.$$

We first need to bound the cardinality of the set $C(\mathbf{x})$.

Lemma B.1 *With probability $1 - c/(\log(d))^2 - N^2 d^{-C}$ we have*

$$|C(\mathbf{x})| \leq C d^{\frac{2}{3}} \log(d).$$

The proof of this lemma is given in Section B.1.

Now, for a vector δ we define:

$$\Delta A(\mathbf{x}, \delta) = \left\{ \ell : \ell \in A(\mathbf{x}) \text{ and } \langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle < 0 \right\},$$

$$\Delta B(\mathbf{x}, \delta) = \left\{ \ell : \ell \in B(\mathbf{x}) \text{ and } \langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle > 0 \right\}.$$

In other words, the set $\Delta A(\mathbf{x}, \delta)$ (respectively $\Delta B(\mathbf{x}, \delta)$) contains all the indices ℓ in $A(\mathbf{x})$ (respectively $B(\mathbf{x})$) in which the sign of $\langle \mathbf{w}_\ell, \mathbf{x} \rangle$ and $\langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle$ are different. We now prove that when $\|\delta\|_{\ell_2} \leq \varepsilon$ we have

$$(B.4) \quad |\Delta A(\mathbf{x}, \delta)|, |\Delta B(\mathbf{x}, \delta)| \leq C d^{\frac{2}{3}},$$

for an absolute constant $C > 0$. By definition, on the event $\mathcal{E}_{\mathbf{W}}$, $\mathbf{W}$ has bounded operator norm, say at most C for some constant $C > 0$. In addition, $\|\delta\|_{\ell_2} \leq \varepsilon$. Therefore, $\|\mathbf{W}\delta\|_{\ell_2} \leq \varepsilon C$. As a result, the number of entries of the vector $\mathbf{W}\delta$ whose absolute value is larger than $d^{-\frac{1}{3}}$ is bounded by $\varepsilon^2 C^2 d^{\frac{2}{3}}$. But from the definitions of the sets $A(\mathbf{x})$ and $\Delta A(\mathbf{x}, \delta)$ it is immediate that for $\ell \in \Delta A(\mathbf{x}, \delta)$ we have $|\langle \mathbf{w}_\ell, \delta \rangle| > d^{-\frac{1}{3}}$. And this results in the fact that $|\Delta A(\mathbf{x}, \delta)| \leq C' d^{\frac{2}{3}}$ with $C' = \varepsilon^2 C^2$. The same argument holds for $|\Delta B(\mathbf{x}, \delta)|$.

Let us now consider the two vectors $\sigma(\mathbf{W}\mathbf{x})$ and $\sigma(\mathbf{W}(\mathbf{x} + \delta))$. We would like to find out how these vectors are different on the indices that belong to the set $A(\mathbf{x})$ or $B(\mathbf{x})$. Let us start with the indices in $B(\mathbf{x})$. Note that for any $\ell \in B(\mathbf{x})$ we have $\langle \mathbf{w}_\ell, \mathbf{x} \rangle < 0$. For this entry, it is easy to see that the two vectors $\sigma(\mathbf{W}\mathbf{x})$ and $\sigma(\mathbf{W}(\mathbf{x} + \delta))$ take different values only if we also have $\ell \in \Delta B(\mathbf{x}, \delta)$. As a result, we can conclude that the two vectors $\sigma(\mathbf{W}\mathbf{x})$ and $\sigma(\mathbf{W}(\mathbf{x} + \delta))$ are the same on all the indices belonging to the set $B(\mathbf{x})$ except at most $C d^{\frac{2}{3}}$ indices. In other words, the difference $\sigma(\mathbf{W}(\mathbf{x} + \delta)) - \sigma(\mathbf{W}\mathbf{x})$ takes zero value on all the indices belonging to the set $B(\mathbf{x})$ except at most $C d^{\frac{2}{3}}$ indices.

For the indices in the set $A(\mathbf{x})$ the situation is different as we are operating in the non-constant part of the ReLU function (note that for any $\ell \in A(\mathbf{x})$ we have $\langle \mathbf{w}_\ell, \mathbf{x} \rangle > 0$). We first claim the following: The two vectors $\sigma(\mathbf{W}(\mathbf{x} + \delta))$ and $\mathbf{W}(\mathbf{x} + \delta) - 1/\sqrt{2\pi}$ are the same on all the entries in the set $A(\mathbf{x})$ except the indices in the set $\Delta A(\mathbf{x}, \delta)$. The justification is as follows: Consider an index $\ell \in A(\mathbf{x})$ such that $\sigma(\langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle) \neq \langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle - 1/\sqrt{2\pi}$. Since $\ell \in A(\mathbf{x})$, we have $\sigma(\langle \mathbf{w}_\ell, \mathbf{x} \rangle) = \langle \mathbf{w}_\ell, \mathbf{x} \rangle - 1/\sqrt{2\pi}$. Now, since $\sigma(\langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle) \neq \langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle - 1/\sqrt{2\pi}$ only if $\langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle < 0$, we obtain that $\sigma(\langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle) \neq \langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle - 1/\sqrt{2\pi}$ only if $\ell \in \Delta A(\mathbf{x}, \delta)$.

In summary, we have shown that (i) On indices belonging to the set $A(\mathbf{x}) \setminus \Delta A(\mathbf{x}, \delta)$: the two vectors $\sigma(\mathbf{W}(\mathbf{x} + \delta))$ and $\mathbf{W}(\mathbf{x} + \delta) - 1/\sqrt{2\pi}$ are the same, and (ii) on the indices belonging to the set $B(\mathbf{x}) \setminus \Delta B(\mathbf{x}, \delta)$ the vector $\sigma(\mathbf{W}(\mathbf{x} + \delta)) - \sigma(\mathbf{W}\mathbf{x})$ takes value 0. Also, (iii) both sets $\Delta A(\mathbf{x}, \delta)$ and $\Delta B(\mathbf{x}, \delta)$ have cardinality at most $C d^{\frac{2}{3}}$. We can thus write:

$$(B.5) \quad \begin{aligned} & \boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x} + \delta)) - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}) \\ &= \sum_{\ell \in A(\mathbf{x})} \theta_\ell (\sigma(\langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle) - \sigma(\langle \mathbf{w}_\ell, \mathbf{x} \rangle)) + \sum_{\ell \in B(\mathbf{x})} \theta_\ell (\sigma(\langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle) - \sigma(\langle \mathbf{w}_\ell, \mathbf{x} \rangle)) \\ & \quad + \sum_{\ell \in C(\mathbf{x})} \theta_\ell (\sigma(\langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle) - \sigma(\langle \mathbf{w}_\ell, \mathbf{x} \rangle)). \end{aligned}$$

We first bound the second and third terms. Using the fact that $|\sigma(a+b) - \sigma(b)| \leq |a|$ we obtain for the third term that:

$$\left| \sum_{\ell \in C(\mathbf{x})} \theta_\ell (\sigma(\langle \mathbf{w}_\ell, \mathbf{x} + \delta \rangle) - \sigma(\langle \mathbf{w}_\ell, \mathbf{x} \rangle)) \right| \leq \sum_{\ell \in C(\mathbf{x})} |\theta_\ell| |\langle \mathbf{w}_\ell, \delta \rangle|$$

$$\begin{aligned}
&\leq \|\boldsymbol{\theta}\|_\infty \sqrt{|C(\mathbf{x})|} \|\mathbf{W}\boldsymbol{\delta}\|_{\ell_2} \\
&\leq \frac{C}{d^{\frac{1}{2}}} (Cd^{\frac{2}{3}} \log(d))^{\frac{1}{2}} C\varepsilon \\
&\leq C' d^{-\frac{1}{6}} \log(d),
\end{aligned}
\tag{B.6}$$

where we have use the result of Lemma B.1 to bound the size of the set $C(\mathbf{x})$ (and hence the above holds with the probability given in that lemma). For the second term we have

$$\left| \sum_{\ell \in B(\mathbf{x})} \theta_\ell (\sigma(\langle \mathbf{w}_\ell, \mathbf{x} + \boldsymbol{\delta} \rangle) - \sigma(\langle \mathbf{w}_\ell, \mathbf{x} \rangle)) \right| \leq \sum_{\ell \in \Delta B(\mathbf{x}, \boldsymbol{\delta})} |\theta_\ell| |\langle \mathbf{w}_\ell, \boldsymbol{\delta} \rangle| \leq \|\boldsymbol{\theta}\|_\infty \sqrt{|\Delta B(\mathbf{x}, \boldsymbol{\delta})|} \|\mathbf{W}\boldsymbol{\delta}\|_{\ell_2} \leq C' d^{-\frac{1}{6}}.
\tag{B.7}$$

Finally, in a similar manner as above we can bound

$$\left| \sum_{\ell \in \Delta A(\mathbf{x}, \boldsymbol{\delta})} \theta_\ell (\sigma(\langle \mathbf{w}_\ell, \mathbf{x} + \boldsymbol{\delta} \rangle) - \sigma(\langle \mathbf{w}_\ell, \mathbf{x} \rangle)) \right| \leq \|\boldsymbol{\theta}\|_\infty \sqrt{|\Delta A(\mathbf{x}, \boldsymbol{\delta})|} \|\mathbf{W}\boldsymbol{\delta}\|_{\ell_2} \leq C' d^{-\frac{1}{6}}.
\tag{B.8}$$

By plugging (B.6), (B.7), and (B.8) into (B.5) we have shown that

$$\begin{aligned}
\boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x} + \boldsymbol{\delta})) - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}) &= \sum_{\ell \in A(\mathbf{x}) \setminus \Delta A(\mathbf{x}, \boldsymbol{\delta})} \theta_\ell (\sigma(\langle \mathbf{w}_\ell, \mathbf{x} + \boldsymbol{\delta} \rangle) - \sigma(\langle \mathbf{w}_\ell, \mathbf{x} \rangle)) + C' d^{-\frac{1}{6}} \log(d) \\
&= \sum_{\ell \in A(\mathbf{x}) \setminus \Delta A(\mathbf{x}, \boldsymbol{\delta})} \theta_\ell (\langle \mathbf{w}_\ell, \mathbf{x} + \boldsymbol{\delta} \rangle - \langle \mathbf{w}_\ell, \mathbf{x} \rangle) + C' d^{-\frac{1}{6}} \log(d) \\
&= \sum_{\ell \in A(\mathbf{x}) \setminus \Delta A(\mathbf{x}, \boldsymbol{\delta})} \theta_\ell \langle \mathbf{w}_\ell, \boldsymbol{\delta} \rangle + C' d^{-\frac{1}{6}} \log(d),
\end{aligned}
\tag{B.9}$$

where the second equality follows from the definition of the set $\Delta A(\mathbf{x}, \boldsymbol{\delta})$.

As a final step, we define the set $A^+(\mathbf{x}) = \{\ell : \langle \mathbf{w}_\ell, \mathbf{x} \rangle > 0\}$. Note that $A(\mathbf{x}) \subseteq A^+(\mathbf{x})$ and $A^+(\mathbf{x}) \setminus A(\mathbf{x}) \subseteq C(\mathbf{x})$. As a result $|A^+(\mathbf{x}) \setminus A(\mathbf{x})| \leq Cd^{\frac{2}{3}} \log(d)$. We thus obtain

$$\begin{aligned}
\sum_{\ell \in A(\mathbf{x}) \setminus \Delta A(\mathbf{x}, \boldsymbol{\delta})} \theta_\ell \langle \mathbf{w}_\ell, \boldsymbol{\delta} \rangle &= \sum_{\ell \in A^+(\mathbf{x})} \theta_\ell \langle \mathbf{w}_\ell, \boldsymbol{\delta} \rangle - \sum_{\ell \in (A^+(\mathbf{x}) \setminus A(\mathbf{x})) \cup \Delta A(\mathbf{x}, \boldsymbol{\delta})} \theta_\ell \langle \mathbf{w}_\ell, \boldsymbol{\delta} \rangle \\
&= \sum_{\ell \in A^+(\mathbf{x})} \theta_\ell \langle \mathbf{w}_\ell, \boldsymbol{\delta} \rangle + C' d^{-\frac{1}{6}} \log(d),
\end{aligned}
\tag{B.10}$$

where the last relations follows from the fact that $|(A^+(\mathbf{x}) \setminus A(\mathbf{x})) \cup \Delta A(\mathbf{x}, \boldsymbol{\delta})| \leq |C(\mathbf{x})| + |\Delta A(\mathbf{x}, \boldsymbol{\delta})| = O(d^{\frac{2}{3}} \log(d))$. By plugging (B.10) into (B.9) we have

$$\begin{aligned}
\boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x} + \boldsymbol{\delta})) - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}) &= \sum_{\ell \in A^+(\mathbf{x})} \theta_\ell \langle \mathbf{w}_\ell, \boldsymbol{\delta} \rangle + C' d^{-\frac{1}{6}} \log(d) \\
&= \sum_{\ell : \langle \mathbf{w}_\ell, \mathbf{x} \rangle > 0} \theta_\ell \langle \mathbf{w}_\ell, \boldsymbol{\delta} \rangle + C' d^{-\frac{1}{6}} \log(d) \\
&= \boldsymbol{\theta}^\top \text{diag}(\mathbf{W}\mathbf{x} > 0) \mathbf{W}\boldsymbol{\delta} + C' d^{-\frac{1}{6}} \log(d),
\end{aligned}$$

which is the result of (B.1).

B.1. Proof of Lemma B.1 Define the random variables $\mu_\ell := \mathbf{w}_\ell^\top \mathbf{x}$ for $\ell = 1, \dots, N$. Note that since $\mathbf{x}$ is gaussian and $\|\mathbf{w}_\ell\|_{\ell_2} = 1$, then $\mu_\ell \sim \mathcal{N}(0, 1)$. Also, note that μ_ℓ 's are correlated with each other and each pair (μ_ℓ, μ_k) is a jointly-gaussian random variable with correlation $\rho_{\ell, k} := \mathbb{E}[\mu_\ell, \mu_k] = \mathbf{w}_\ell^\top \mathbf{w}_k$. Define $z_\ell = \mathbf{1}\{\mu_\ell \in [-d^{-\frac{1}{3}}, d^{-\frac{1}{3}}]\}$. Note that $\mathbb{E}[z_\ell] \leq cd^{-\frac{1}{3}}$.

Define the event $\mathcal{E} := \{|\mathbf{w}_\ell^\top \mathbf{w}_k| \leq d^{-1/2} \sqrt{C \log(d)}, \forall \ell, k \in [N]\}$. Since $\mathbf{w}_\ell \sim_{i.i.d} \text{Unif}(\mathbb{S}^{d-1})$, it is easy to see that $\mathbb{P}(\mathcal{E}) \geq 1 - N^2 d^{-C}$. We also have

$$\mathbb{P}(\mu_\ell = 1, \mu_k = 1) = \int_{\mu \in [-d^{-\frac{1}{3}}, d^{-\frac{1}{3}}]} f(\mu_\ell = \mu) \mathbb{P}(\mu_k \in [-d^{-\frac{1}{3}}, d^{-\frac{1}{3}}] \mid \mu_\ell = \mu) d\mu,$$

where f denotes pdf of μ_ℓ . Now, given $\mu_\ell = \mu$, the distribution of μ_k is $\mathcal{N}(\mu \rho_{\ell,k}, (1 - \rho_{\ell,k}^2))$. It is easy to see that on the event $\mathcal{E}$, for $|\mu| \leq d^{-\frac{1}{3}}$ we have

$$\mathbb{P}(\mu_k \in [-d^{-\frac{1}{3}}, d^{-\frac{1}{3}}] \mid \mu_\ell = \mu) \leq cd^{-\frac{1}{3}},$$

and thus

$$\mathbb{E}[z_\ell z_k] = \mathbb{P}(\mu_\ell = 1, \mu_k = 1) \leq cd^{-\frac{1}{3}} \int_{\mu \in [-d^{-\frac{1}{3}}, d^{-\frac{1}{3}}]} f(\mu_\ell = \mu) d\mu \leq cd^{-\frac{2}{3}}.$$

Let us now consider the average $\bar{z} = \frac{1}{N} \sum_{\ell=1}^N z_\ell$. We have

$$\mathbb{E}[|\bar{z} - \mathbb{E}[\bar{z}]|^2] = \mathbb{E}[\bar{z}^2] - \mathbb{E}[\bar{z}]^2 \leq \mathbb{E}[\bar{z}^2] \leq \frac{1}{N^2} \sum_{\ell,k=1}^N \mathbb{E}[z_\ell z_k] \leq cd^{-\frac{2}{3}},$$

and thus we obtain via the Chebyshev's inequality that

$$\mathbb{P}\left\{|\bar{z} - \mathbb{E}[\bar{z}]| \geq d^{-\frac{1}{3}} \log(d); \mathcal{E}\right\} \leq \frac{c}{(\log(d))^2}.$$

Now, by noticing that $|C(\mathbf{x})|/N = \bar{z}$, and $\mathbb{E}[\bar{z}] \leq cd^{-\frac{1}{3}}$, along with the assumption that N, d grows proportionally we obtain

$$\mathbb{P}\left\{|C(\mathbf{x})| \geq Cd \times d^{-\frac{1}{3}} \log(d); \mathcal{E}\right\} \leq \frac{c}{(\log(d))^2}.$$

Finally, we have

$$\mathbb{P}\left\{|C(\mathbf{x})| \geq Cd^{\frac{2}{3}} \log(d)\right\} \leq \frac{c}{\log(d)^2} + \mathbb{P}(\mathcal{E}^c) \leq \frac{c}{(\log(d))^2} + N^2 d^{-C}.$$

B.2. Proof of Proposition 6.2

Recall the loss $\mathcal{L}(\boldsymbol{\theta})$ given by

$$\mathcal{L}(\boldsymbol{\theta}) := \max_{\|\boldsymbol{\delta}_i\|_{\ell_2} \leq \varepsilon} \frac{1}{2n} \sum_{i=1}^n (y_i - \boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i)))^2 + \frac{\zeta}{2} \boldsymbol{\theta}^\top \boldsymbol{\Omega} \boldsymbol{\theta},$$

and $\widehat{\boldsymbol{\theta}} = \arg \min \mathcal{L}(\boldsymbol{\theta})$. Bounding $\|\widehat{\boldsymbol{\theta}}\|_{\ell_2}$ is straightforward. By optimality of $\widehat{\boldsymbol{\theta}}$ and comparing the loss at $\widehat{\boldsymbol{\theta}}$ and $\mathbf{0}$ we get

$$\frac{\zeta}{2} \widehat{\boldsymbol{\theta}}^\top \boldsymbol{\Omega} \widehat{\boldsymbol{\theta}} \leq \mathcal{L}(\mathbf{0}) = \frac{1}{2n} \sum_{i=1}^n y_i^2 < C,$$

with probability at least $1 - e^{-cn}$, for absolute constants $c, C > 0$. Since $\boldsymbol{\Omega} \geq \mathbf{I}$, this implies that $\|\boldsymbol{\theta}\|_{\ell_2} \leq C_0$ for sufficiently large C_0 .

To bound $\|\widehat{\boldsymbol{\theta}}\|_{\ell_\infty}$ we just need to bound any given entry of $\widehat{\boldsymbol{\theta}}$, e.g. its last entry, with high probability. By symmetry, all the entries have the same marginal distribution. Consequently, each entry of $\widehat{\boldsymbol{\theta}}$ can be analyzed in the same way and $\|\widehat{\boldsymbol{\theta}}\|_{\ell_\infty}$ can then be controlled by using the union bound.

With a slight abuse of notation, we consider a $(N + 1)$ dimensional version of the above optimization over $[\boldsymbol{\theta}; u]$ and denote the last coordinate of the optimal solution by $\hat{u}$. Let $\lambda := \frac{\sqrt{\log(d)}}{d}$ and

$$\boldsymbol{\Omega} = \begin{pmatrix} \tilde{\boldsymbol{\Omega}} & \lambda \mathbf{1} \\ \lambda \mathbf{1}^\top & \lambda + 1 \end{pmatrix},$$

where $\tilde{\boldsymbol{\Omega}}$ is of size N and $\mathbf{1} = (1, 1, \dots, 1) \in \mathbb{R}^N$. The last coordinate $\hat{u}$ can be expressed as

$$\begin{aligned} \hat{u} = \arg \min_u \min_{\boldsymbol{\theta}} & \left[\frac{1}{2n} \sum_{i=1}^n \max_{\|\boldsymbol{\delta}_i\|_{\ell_2} \leq \varepsilon} \left(y_i - \boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i^*)) - u \sigma(\mathbf{w}_{N+1}^\top (\mathbf{x}_i + \boldsymbol{\delta}_i^*)) \right)^2 \right. \\ & \left. + \frac{\zeta}{2} (\boldsymbol{\theta}^\top \tilde{\boldsymbol{\Omega}} \boldsymbol{\theta} + 2\lambda(\mathbf{1}^\top \boldsymbol{\theta})u + (\lambda + 1)u^2) \right]. \end{aligned} \quad (\text{B.11})$$

We next define $f(u)$ as the objective function of u in (B.11). We proceed by deriving a lower bound for $f(u)$.

Let $\boldsymbol{\theta}_*$ be the optimal $\boldsymbol{\theta}$ if we set $u = 0$ and denote by $\boldsymbol{\delta}_i^{\setminus u}$ the maximizing $\boldsymbol{\delta}_i$, when $u = 0$. Note that $\boldsymbol{\delta}_i^{\setminus u}$ is in general a function of $\boldsymbol{\theta}$. In addition, define

$$\begin{aligned} \ell([\boldsymbol{\theta}, u]) &= \frac{1}{2n} \sum_{i=1}^n \left(y_i - \boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u})) - u \sigma(\mathbf{w}_{N+1}^\top (\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u})) \right)^2, \\ Q([\boldsymbol{\theta}, u]) &= \frac{\zeta}{2} (\boldsymbol{\theta}^\top \tilde{\boldsymbol{\Omega}} \boldsymbol{\theta} + 2\lambda(\mathbf{1}^\top \boldsymbol{\theta})u + (\lambda + 1)u^2). \end{aligned}$$

Since the pointwise maximum of convex functions is convex, the function $\ell(\cdot)$ is convex and hence we have

$$(\text{B.12}) \quad \ell([\boldsymbol{\theta}; u]) \geq \ell([\boldsymbol{\theta}_*; 0]) + \langle \nabla_{\boldsymbol{\theta}} \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*; 0]}, \boldsymbol{\theta} - \boldsymbol{\theta}_* \rangle + \nabla_u \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*; 0]} u.$$

For quadratic function $Q([\boldsymbol{\theta}; u])$ we have

$$\begin{aligned} Q([\boldsymbol{\theta}; u]) &= \frac{\zeta}{2} \boldsymbol{\theta}_*^\top \tilde{\boldsymbol{\Omega}} \boldsymbol{\theta}_* + \zeta \boldsymbol{\theta}_*^\top \tilde{\boldsymbol{\Omega}} (\boldsymbol{\theta} - \boldsymbol{\theta}_*) + \frac{\zeta}{2} (\boldsymbol{\theta} - \boldsymbol{\theta}_*)^\top \tilde{\boldsymbol{\Omega}} (\boldsymbol{\theta} - \boldsymbol{\theta}_*) + \frac{\zeta}{2} (2\lambda u \mathbf{1}^\top \boldsymbol{\theta} + (\lambda + 1)u^2) \\ (\text{B.13}) \quad &= Q([\boldsymbol{\theta}_*; 0]) + \zeta (\boldsymbol{\theta}_*^\top \tilde{\boldsymbol{\Omega}} (\boldsymbol{\theta} - \boldsymbol{\theta}_*) + \lambda(\mathbf{1}^\top \boldsymbol{\theta}_*)u) + \frac{\zeta}{2} \{ (\boldsymbol{\theta} - \boldsymbol{\theta}_*)^\top \tilde{\boldsymbol{\Omega}} (\boldsymbol{\theta} - \boldsymbol{\theta}_*) \\ (\text{B.14}) \quad &+ (\lambda + 1)u^2 + 2\lambda u \mathbf{1}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_*) \}. \end{aligned}$$

Combining (B.12) and (B.13) we get

$$\begin{aligned} \mathcal{L}([\boldsymbol{\theta}; u]) &\geq \ell([\boldsymbol{\theta}; u]) + Q([\boldsymbol{\theta}; u]) \\ &\geq \mathcal{L}([\boldsymbol{\theta}_*; 0]) + \langle \nabla_{\boldsymbol{\theta}} \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*; 0]}, \boldsymbol{\theta} - \boldsymbol{\theta}_* \rangle + \zeta \boldsymbol{\theta}_*^\top \tilde{\boldsymbol{\Omega}} (\boldsymbol{\theta} - \boldsymbol{\theta}_*) + (\nabla_u \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*; 0]} + \zeta \lambda \mathbf{1}^\top \boldsymbol{\theta}_*)u \\ (\text{B.15}) \quad &+ \frac{\zeta}{2} \{ (\boldsymbol{\theta} - \boldsymbol{\theta}_*)^\top \tilde{\boldsymbol{\Omega}} (\boldsymbol{\theta} - \boldsymbol{\theta}_*) + (\lambda + 1)u^2 + 2\lambda u \mathbf{1}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_*) \}. \end{aligned}$$

Here, the first inequality holds since $\mathcal{L}(\cdot)$ involves maximization over $\boldsymbol{\delta}_i$, while in definition of $\ell(\cdot)$ we consider $\boldsymbol{\delta}_i^{\setminus u}$. Though, note that $\mathcal{L}([\boldsymbol{\theta}_*; 0]) = \ell([\boldsymbol{\theta}_*; 0]) + Q([\boldsymbol{\theta}_*; 0])$ because when $u = 0$, $\boldsymbol{\delta}_i^{\setminus u}$ are the maximizing perturbations by definition. We used this observation in the second inequality above.

We argue that the second term in the right-hand side is zero. To see this, first write the partial derivative $\nabla_{\boldsymbol{\theta}} \ell$ as

$$\begin{aligned} \nabla_{\boldsymbol{\theta}} \ell([\boldsymbol{\theta}, u]) &= -\frac{1}{n} \sum_{i=1}^n \left(y_i - \boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u})) - u \sigma(\mathbf{w}_{N+1}^\top (\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u})) \right) \times \\ &\quad \left(\frac{\partial}{\partial \boldsymbol{\theta}} [\boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u}))] + u \frac{\partial}{\partial \boldsymbol{\theta}} \sigma(\mathbf{w}_{N+1}^\top (\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u})) \right). \end{aligned}$$

(Note that $\boldsymbol{\delta}_i^{\setminus u}$ is a function of $\boldsymbol{\theta}$.) Therefore,

(B.16)

$$\nabla_{\boldsymbol{\theta}} \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*, 0]} = -\frac{1}{n} \sum_{i=1}^n \left(y_i - \boldsymbol{\theta}_*^\top \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u})) \right) \left(\frac{\partial}{\partial \boldsymbol{\theta}} [\boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u}))] \Big|_{[\boldsymbol{\theta}_*, 0]} \right).$$

By using the first-order optimality condition for $\boldsymbol{\theta}_*$ we have

$$\begin{aligned} \nabla_{\boldsymbol{\theta}} \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*, 0]} + \zeta \boldsymbol{\theta}_*^\top \tilde{\boldsymbol{\Omega}} &= -\frac{1}{n} \sum_{i=1}^n \left(y_i - \boldsymbol{\theta}_*^\top \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u})) \right) \left(\frac{\partial}{\partial \boldsymbol{\theta}} [\boldsymbol{\theta}^\top \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u}))] \Big|_{[\boldsymbol{\theta}_*, 0]} \right) + \zeta \boldsymbol{\theta}_*^\top \tilde{\boldsymbol{\Omega}} = 0. \end{aligned}$$

Using the above relation in (B.15) we arrive at

$$\begin{aligned} \mathcal{L}([\boldsymbol{\theta}; u]) &\geq \mathcal{L}([\boldsymbol{\theta}_*; 0]) + (\nabla_u \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*, 0]} + \zeta \lambda \mathbf{1}^\top \boldsymbol{\theta}_*) u \\ &\quad + \frac{\zeta}{2} \{ (\boldsymbol{\theta} - \boldsymbol{\theta}_*)^\top \tilde{\boldsymbol{\Omega}} (\boldsymbol{\theta} - \boldsymbol{\theta}_*) + (\lambda + 1) u^2 + 2\lambda u \mathbf{1}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_*) \}. \end{aligned}$$

Therefore, by minimizing the both sides over $\boldsymbol{\theta}$ we obtain

$$\begin{aligned} f(u) &= \min_{\boldsymbol{\theta}} \mathcal{L}([\boldsymbol{\theta}; u]) \\ &\geq f(0) + (\nabla_u \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*, 0]} + \zeta \lambda \mathbf{1}^\top \boldsymbol{\theta}_*) u \\ &\quad + \min_{\boldsymbol{\theta}} \frac{\zeta}{2} \{ (\boldsymbol{\theta} - \boldsymbol{\theta}_*)^\top \tilde{\boldsymbol{\Omega}} (\boldsymbol{\theta} - \boldsymbol{\theta}_*) + (\lambda + 1) u^2 + 2\lambda u \mathbf{1}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_*) \} \\ (B.17) \quad &= f(0) + (\nabla_u \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*, 0]} + \lambda \mathbf{1}^\top \boldsymbol{\theta}_*) u + \frac{\zeta}{2} u^2 (1 + \lambda - \lambda^2 \mathbf{1}^\top \tilde{\boldsymbol{\Omega}}^{-1} \mathbf{1}). \end{aligned}$$

By definition of $\tilde{\boldsymbol{\Omega}}$, it has $\mathbf{1}$ as an eigenvector with eigenvalue $1 + \lambda d$. So,

$$(B.18) \quad 1 + \lambda - \lambda^2 \mathbf{1}^\top \tilde{\boldsymbol{\Omega}}^{-1} \mathbf{1} \geq 1 + \lambda - \frac{\lambda^2 d}{1 + \lambda d} > 1.$$

By optimality of $\hat{u}$, we have $f(\hat{u}) \leq f(0)$, which together with (B.17) and (B.18) imply that

$$(B.19) \quad |\hat{u}| \leq \frac{2}{\zeta} \left| \nabla_u \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*, 0]} + \zeta \lambda \mathbf{1}^\top \boldsymbol{\theta}_* \right|.$$

We next bound the terms on the right-hand side separately. We have

$$\lambda \mathbf{1}^\top \boldsymbol{\theta}_* \leq \lambda \|\mathbf{1}\|_{\ell_2} \|\boldsymbol{\theta}_*\|_{\ell_2} \leq \sqrt{\frac{\log(d)}{d}} \|\boldsymbol{\theta}_*\|_{\ell_2}.$$

By optimality of $\boldsymbol{\theta}_*$ (when we set $u = 0$) and comparing it with $\mathbf{0}$ we get

$$\frac{\zeta}{2} \boldsymbol{\theta}_*^\top \tilde{\boldsymbol{\Omega}} \boldsymbol{\theta}_* \leq \frac{1}{2n} \sum_{i=1}^n y_i^2 < C,$$

with probability at least $1 - e^{-cn}$.

Sine $\tilde{\Omega} \geq \mathbf{I}$, this implies that $\lambda \mathbf{1}^\top \boldsymbol{\theta}_* = O_{\mathbb{P}}(\sqrt{\log(d)/d})$.

To bound the other term, recall that by definition $\frac{\partial}{\partial u} \boldsymbol{\delta}_i^{\setminus u} = 0$ and so $\nabla_u \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*, 0]}$ is given by

$$(B.20) \quad \nabla_u \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*, 0]} = \frac{1}{n} \sum_{i=1}^n \left(y_i - \boldsymbol{\theta}_*^\top \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u})) \right) \sigma(\mathbf{w}_{N+1}^\top (\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u})).$$

To simplify the notation define $m_i := \frac{1}{\sqrt{n}} (y_i - \boldsymbol{\theta}_*^\top \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u})))$ and $\mathbf{X} = [\mathbf{x}_1 | \dots | \mathbf{x}_n]^\top$. Consider the following event:

$$\mathcal{E} := \left\{ \|\mathbf{m}\|_{\ell_2} \leq C, \frac{1}{\sqrt{d}} \|\mathbf{X}\| \leq C \right\},$$

where $\mathbf{m} = (m_1, \dots, m_n)^\top$ and $C > 0$ is a sufficiently large constant. We show that $\mathcal{E}$ is a high probability event. To see this, first observe that

$$(B.21) \quad \|\mathbf{m}\|_{\ell_2}^2 = \frac{1}{n} \sum_{i=1}^n \left(y_i - \boldsymbol{\theta}_*^\top \sigma(\mathbf{W}(\mathbf{x}_i + \boldsymbol{\delta}_i^{\setminus u})) \right)^2 \leq \frac{1}{n} \sum_{i=1}^n y_i^2,$$

where the inequality follows from optimality of $\boldsymbol{\theta}_*$ and comparing the loss value $\mathcal{L}([\boldsymbol{\theta}_*, 0])$ with $\mathcal{L}([0; 0])$. Therefore,

$$\mathbb{P}(\|\mathbf{m}\|_{\ell_2} > C) \leq \mathbb{P}\left(\frac{1}{\sqrt{n}} \|\mathbf{y}\|_{\ell_2} > C\right) \leq e^{-C'n},$$

for absolute constants C, C' (depending on the noise variance τ^2). Also, given that $\mathbf{X}$ has i.i.d standard normal entries we have

$$\mathbb{P}\left(\frac{1}{\sqrt{d}} \|\mathbf{X}\| > C\right) \leq 2e^{-cn}.$$

Putting the last two bounds together we obtain $\mathbb{P}(\mathcal{E}^c) \leq 3e^{-cn}$.

Let $\mathcal{F}$ be the σ -algebra generated by an arbitrary $\mathbf{X}, \mathbf{W}, \mathbf{y}$ in $\mathcal{E}$. Clearly, m_i are measurable with respect to $\mathcal{F}$. Also, $\mathbf{w}_{N+1}$ is drawn independently from $\mathcal{F}$ and hence conditioned on that $\mathbf{w}_{N+1}^\top \mathbf{x}_i \sim \mathcal{N}(0, 1)$. Since $\mathbb{E}[\sigma(G)] = 0$ for $G \sim \mathcal{N}(0, 1)$, we have $\mathbb{E}[\sigma(\mathbf{w}_{N+1}^\top \mathbf{x}_i) | \mathcal{F}] = 0$. To bound $\nabla_u \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*, 0]}$ we view that as a function of $\mathbf{w}_{N+1}$ and condition on $\mathcal{F}$. We then have

$$\mathbb{E}[\nabla_u \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*, 0]} | \mathcal{F}] = \frac{1}{\sqrt{n}} \sum_{i=1}^n m_i \mathbb{E}[\sigma(\mathbf{w}_{N+1}^\top \mathbf{x}_i) | \mathcal{F}] = 0.$$

Also this is a Lipschitz continuous function of $\mathbf{w}_{N+1}$ with a Lipschitz factor at most $\frac{C}{\sqrt{d}} \|\mathbf{X} \mathbf{m}\|_{\ell_2} \leq \frac{1}{\sqrt{\psi_2}} \frac{1}{\sqrt{d}} \|\mathbf{X}\| \|\mathbf{m}\|_{\ell_2} \leq \frac{C^2}{\sqrt{\psi_2}} := C_0$. Since $\mathbf{w}_{N+1}$ is chosen uniformly at random from the unit sphere, we can apply the concentration bound for Lipschitz function (see e.g. [92, Theorem 5.1.4]), which implies that

$$(B.22) \quad \mathbb{P}(\nabla_u \ell([\boldsymbol{\theta}, u])|_{[\boldsymbol{\theta}_*, 0]} > t) \leq 2e^{-c't^2}.$$

Choosing $t = C\sqrt{\frac{\log(d)}{d}}$ and using this bound in (B.19) we get

$$|\hat{u}| \leq C' \sqrt{\frac{\log(d)}{d}},$$

with probability at least $1 - 4e^{-cn} - 2d^{-c'C^2}$. The result follows by choosing $C > 0$ large enough so that $c'C^2 > 1$ and union bounding over the N coordinates of $\hat{\boldsymbol{\theta}}$, along with the assumption that N, n, d grow at the same order.

B.3. Proof of Proposition 6.3 We know from the result of Proposition 6.1, or more precisely the equations (B.1)-(B.2) in its proof, that with probability $1 - o_d(1)$ we have

$$\sup_{\theta \in \mathcal{C}_\theta} \left| \max_{\|\delta\|_{\ell_2} \leq \varepsilon} |y - \theta^\top \sigma(\mathbf{W}(x + \delta))| - (|y - \theta^\top \sigma(\mathbf{W}x)| + \varepsilon \|\mathbf{W}^\top \text{diag}(\mathbb{I}(\mathbf{W}x > 0)) \theta\|_{\ell_2}) \right| = \alpha_d,$$

where $\alpha_d = O\left(\frac{\log(d)}{d^{1/6}}\right)$.

One can thus write from (6.1) and (6.5) that for any $\theta \in \mathcal{C}_\theta$

$$\begin{aligned} \left| \mathcal{L}(\theta) - \mathring{\mathcal{L}}(\theta) \right| &\leq \alpha_d^2 + \alpha_d \frac{1}{n} \sum_{i=1}^n \max_{\|\delta_i\|_{\ell_2} \leq \varepsilon} |y_i - \theta^\top \sigma(\mathbf{W}(x_i + \delta_i))| \\ &\leq \alpha_d^2 + \alpha_d \frac{1}{n} \sum_{i=1}^n |y_i - \theta^\top \sigma(\mathbf{W}x_i)| + \alpha_d \|\theta\|_{\ell_2} \|\mathbf{W}\| \varepsilon \\ &\leq \alpha_d^2 + \alpha_d \frac{1}{n} \sum_{i=1}^n |y_i - \theta^\top \sigma(\mathbf{W}x_i)| + c_1 \alpha_d, \end{aligned}$$

where $c_1 > 0$ is an absolute constant. The second inequality follows from the fact that the ReLU function is 1-Lipschitz, and the third inequality follows from $\theta \in \mathcal{C}_\theta$ as well as the fact that $\|\mathbf{W}\|$ is bounded.

We will now show that

$$(B.23) \quad \sup_{\theta \in \mathcal{C}_\theta} \frac{1}{n} \sum_{i=1}^n |y_i - \theta^\top \sigma(\mathbf{W}x_i)| = O_{d, \mathbb{P}}(d^{\frac{1}{12}}).$$

It is easy to see that proving the above relation will finish the proof as $\alpha_d = O\left(\frac{\log(d)}{d^{1/6}}\right)$.

Fix a θ such that $\|\theta\|_{\ell_2} \leq C$. Recall that (x_i, y_i) are generated i.i.d. according to the the distribution (2.1). Since the random variables $|y_i - \theta^\top \sigma(\mathbf{W}x_i)|$ are sub-gaussian (see e.g. (D.93) in Lemma D.9), we can write

$$(B.24) \quad \mathbb{P}\left(\frac{1}{n} \sum_{i=1}^n |y_i - \theta^\top \sigma(\mathbf{W}x_i)| \geq d^{\frac{1}{12}}\right) \leq c_2 e^{-c_2 d^{\frac{7}{6}}},$$

for an absolute constant $c_2 > 0$.

Now, to prove (B.23), we use an ϵ -net argument. Consider a 1-net of the set $\{\theta : \|\theta\|_{\ell_2} \leq C\}$. We know that such a 1-net $\mathcal{S}$ exists with size at most $|\mathcal{S}| \leq 2^{c_3 d}$ where $c_3 > 0$ is an absolute constant. Let $\theta_1 \in \mathcal{S}$ be a vector in this net, and consider another vector θ_2 in the 1-neighborhood of θ_1 – i.e. $\|\theta_1 - \theta_2\|_{\ell_2} \leq 1$. We can write

$$\begin{aligned} \left| \frac{1}{n} \sum_{i=1}^n |y_i - \theta_1^\top \sigma(\mathbf{W}x_i)| - \frac{1}{n} \sum_{i=1}^n |y_i - \theta_2^\top \sigma(\mathbf{W}x_i)| \right| &\leq \frac{1}{n} \sum_{i=1}^n |(\theta_1 - \theta_2)^\top \sigma(\mathbf{W}x_i)| \\ &= \frac{1}{n} \|(\theta_1 - \theta_2)^\top \mathbf{M}\|_{\ell_1} \\ &\leq \frac{1}{\sqrt{n}} \|\mathbf{M}\| \|\theta_1 - \theta_2\|_{\ell_2} \\ &\leq \frac{1}{\sqrt{n}} \|\mathbf{M}\|, \end{aligned} \tag{B.25}$$

where the matrix $\mathbf{M}$ is defined as $\mathbf{M} = [\sigma(\mathbf{W}x_1) | \sigma(\mathbf{W}x_2) | \dots | \sigma(\mathbf{W}x_n)]$. Now, since the random vectors $\sigma(\mathbf{W}x_i)$ are independently generated and sub-gaussian (see (D.93)), we can conclude that

$$(B.26) \quad \mathbb{P}(\|\mathbf{M}\| \geq c_4 \sqrt{n}) \leq c_5 e^{-c_5 d},$$

for absolute constants $c_4, c_5 > 0$ (recall that d, n , and N grow proportionally as per Assumption 1). As a result, from (B.25) and (B.26), we have

(B.27)

$$\mathbb{P} \left(\sup_{\theta_1, \theta_2: \|\theta_1 - \theta_2\|_{\ell_2} \leq 1} \left| \frac{1}{n} \sum_{i=1}^n |y_i - \theta_1^\top \sigma(\mathbf{W} \mathbf{x}_i)| - \frac{1}{n} \sum_{i=1}^n |y_i - \theta_2^\top \sigma(\mathbf{W} \mathbf{x}_i)| \right| \geq c_4 \right) \leq c_5 e^{-c_5 d}.$$

Now, by using (B.24) and (B.27), and a union bound argument over $\mathcal{S}$, we obtain:

$$\begin{aligned} \mathbb{P} \left(\sup_{\theta: \|\theta\|_{\ell_2} \leq C} \frac{1}{n} \sum_{i=1}^n |y_i - \theta^\top \sigma(\mathbf{W} \mathbf{x}_i)| \geq d^{\frac{1}{12}} + c_4 \right) &\leq \mathbb{P} \left(\sup_{\theta \in \mathcal{S}} \frac{1}{n} \sum_{i=1}^n |y_i - \theta^\top \sigma(\mathbf{W} \mathbf{x}_i)| \geq d^{\frac{1}{12}} \right) + c_5 e^{-c_5 d} \\ &\leq c_2 e^{c_3 d - c_2 d^{\frac{7}{6}}} + c_5 e^{-c_5 d} = O(e^{-c_5 d}). \end{aligned}$$

The claim (B.23) now follows because $\mathcal{C}_\theta \subseteq \{\theta: \|\theta\|_{\ell_2} \leq C\}$.

APPENDIX C: PROOFS OF STEP 2: CONCENTRATION OF THE ADVERSARIAL EFFECTS

C.1. Proof of Proposition 6.4 Recall the high probability event $\mathcal{E}_\mathbf{W}$ given by (6.2). We also define the event $\mathcal{E}_\mathbf{x} := \{\|\mathbf{x}_i\|_{\ell_2} \leq \sqrt{5d}, \forall i \in [n]\}$. Since $\mathbf{x}_i \sim \mathcal{N}(0, \mathbf{I}_d)$, $\|\mathbf{x}_i\|_{\ell_2}^2 \sim \chi_d^2$ is a chi-squared distribution with d degrees of freedom. Using chi-squared distribution tail bound (see e.g. [49, lemma 1]) along with a union bound over $i \in [n]$, we obtain $\mathbb{P}(\mathcal{E}_\mathbf{x}) \geq 1 - ne^{-d}$. Since d and n grow proportionally as per Assumption 1, both of the events $\mathcal{E}_\mathbf{W}$ and $\mathcal{E}_\mathbf{x}$ are high probability events, and so it suffices to prove the claim 6.7 on the event $\mathcal{E}_\mathbf{W} \cap \mathcal{E}_\mathbf{x}$.

To prove the proposition, we first state the following lemma which establishes a deviation bound for a fixed $\theta \in \mathcal{C}_\theta$ and fixed $i \in [n]$.

Lemma C.1 *For any fixed $\theta \in \mathcal{C}_\theta$ and fixed $i \in [n]$, the following holds :*

$$\mathbb{P}_{\mathbf{x}_i} \left\{ |\eta_i(\theta)^2 - \mathbb{E}[\eta_i(\theta)^2]| \geq \gamma; \mathcal{E}_\mathbf{W} \cap \mathcal{E}_\mathbf{x} \right\} \leq \frac{c \log^6(d)}{d\gamma^2},$$

for some absolute constant $c > 0$.

Proof of Lemma C.1 is given in Section C.2.

Fix $\theta \in \mathcal{C}_\theta$ and recall our notation $\nu_i(\theta; \gamma) := \mathbb{I}(|\eta_i(\theta)^2 - \mathbb{E}[\eta_i(\theta)^2]| > \gamma)$. Given that $\mathbf{x}_i$ are i.i.d, the random variables $\nu_i \in \{0, 1\}$ are also i.i.d. Bernoulli random variables. Therefore,

$$\begin{aligned} \mathbb{P} \left(\frac{1}{n} \sum_{i=1}^n \nu_i(\theta; \gamma) \geq \frac{1}{\sqrt{\log(d\gamma^2)}} \right) &= \mathbb{P} \left(\sum_{i=1}^n \nu_i(\theta; \gamma) \geq \frac{n}{\sqrt{\log(d\gamma^2)}} \right) \\ &\leq \sum_{\ell=\frac{n}{\sqrt{\log(d\gamma^2)}}}^n \binom{n}{\ell} \mathbb{E}[\nu_1(\theta; \gamma)]^\ell (1 - \mathbb{E}[\nu_1(\theta; \gamma)])^{n-\ell} \\ &\leq \mathbb{E}[\nu_1(\theta; \gamma)]^{\frac{n}{\sqrt{\log(d\gamma^2)}}} \sum_{\ell=\frac{n}{\sqrt{\log(d\gamma^2)}}}^n \binom{n}{\ell} \\ &\leq 2^n \left(\frac{c \log^6(d)}{d\gamma^2} \right)^{\frac{n}{\sqrt{\log(d\gamma^2)}}} \leq (2c \log^6(d))^n e^{-n\sqrt{\log(d\gamma^2)}}, \end{aligned} \tag{C.1}$$

where the last step follows from Lemma C.1 by which $\mathbb{E}[\nu_1(\theta; \gamma)] \leq c \log^6(d)/(d\gamma^2)$ on the event $\mathcal{E}_\mathbf{W} \cap \mathcal{E}_\mathbf{x}$.

Note that the above bound was for a fixed $\theta \in \mathcal{C}_\theta$. In order to prove claim 6.7 we use an ε -net argument. We write

$$\begin{aligned}
 \sup_{\theta \in \mathcal{C}_\theta} \frac{1}{n} \sum_{i=1}^n \nu_i(\theta; \gamma) &\leq \sup_{\|\theta\|_{\ell_2} \leq C_0} \frac{1}{n} \sum_{i=1}^n \nu_i(\theta; \gamma) \\
 &= \sup_{\|\theta\|_{\ell_2} = C_0} \frac{1}{n} \sum_{i=1}^n \nu_i(\theta; \gamma) \\
 &= \sup_{\theta \in \mathbb{S}^{d-1}} \frac{1}{n} \sum_{i=1}^n \nu_i(\theta; \frac{\gamma}{C_0^2}),
 \end{aligned}
 \tag{C.2}$$

where the first step follows from definition of $\mathcal{C}_\theta$; the second step follows from a simple scaling argument, and the third step follows from definition of $\eta_i^2(\theta)$ and $\nu_i(\theta; \gamma)$. We recall that $\mathbb{S}^{d-1}$ denotes the unit $(d-1)$ -dimensional sphere.

Next consider a ε -net $\mathcal{N}$ of $\mathbb{S}^{d-1}$ for $\varepsilon = c_0\gamma$. By [91, Lemma 5.2] we can choose the net $\mathcal{N}$ so that $|\mathcal{N}| \leq (1 + \frac{2}{c_0\gamma})^d$. We use the lemma below to relate the quantity $\nu_i(\theta; \gamma)$ for $\theta \in \mathbb{S}^{d-1}$ to a $\theta \in \mathcal{N}$.

Lemma C.2 *For $\theta \in \mathbb{S}^{d-1}$ choose $\tilde{\theta} \in \mathcal{N}$ which approximates θ as $\|\theta - \tilde{\theta}\|_{\ell_2} \leq c_0\gamma$. On the event $\mathcal{E}_W$ we have the following for all $i \in [n]$:*

$$\nu_i(\theta; \gamma) = 1 \implies \nu_i(\tilde{\theta}; \gamma(1 - 2c_0\sqrt{\psi_{1,d}} - 2c_0C)) = 1.
 \tag{C.3}$$

We refer to Section C.3 for the proof of Lemma C.2.

Continuing from (C.2) and using Lemma C.2 we get

$$\begin{aligned}
 \sup_{\theta \in \mathcal{C}_\theta} \frac{1}{n} \sum_{i=1}^n \nu_i(\theta; \gamma) &\leq \sup_{\theta \in \mathbb{S}^{d-1}} \frac{1}{n} \sum_{i=1}^n \nu_i(\theta; \frac{\gamma}{C_0^2}) \\
 &\leq \sup_{\tilde{\theta} \in \mathcal{N}} \frac{1}{n} \sum_{i=1}^n \nu_i(\tilde{\theta}; \frac{\gamma}{C_0^2}(1 - 2c_0\sqrt{\psi_{1,d}} - 2c_0C)).
 \end{aligned}
 \tag{C.4}$$

Let $\tilde{\gamma} := \frac{\gamma}{C_0^2}(1 - 2c_0\sqrt{\psi_{1,d}} - 2c_0C)$. By choosing the constant c_0 small enough we have $\tilde{\gamma} \geq 0$. Using (C.1) along with union-bounding over the net $\mathcal{N}$ we get

$$\mathbb{P}\left(\sup_{\tilde{\theta} \in \mathcal{N}} \frac{1}{n} \sum_{i=1}^n \nu_i(\tilde{\theta}; \tilde{\gamma}) \geq \frac{1}{\sqrt{\log(d\tilde{\gamma}^2)}}\right) \leq \left(1 + \frac{2}{c_0\gamma}\right)^d (2c \log^6(d))^n e^{-n\sqrt{\log(d\tilde{\gamma}^2)}}.$$

Since n and d grow proportionally and also $\gamma, \tilde{\gamma}$ are of same order, the above event is a high probability event if $\log(1/\gamma) = o(\sqrt{\log(d)})$ or equivalently if $\frac{1}{\gamma} = e^{o(\sqrt{\log(d)})}$. The result follows by combining the above bound with (C.4).

C.2. Proof of Lemma C.1 We decompose the step function as

$$\mathbb{I}(z > 0) = \mu_0 + \mu_1 z + \mu_* \varphi(z),$$

where for $G \sim \mathcal{N}(0, 1)$,

$$\mu_0 := \mathbb{E}[\mathbb{I}(G > 0)] = \frac{1}{2}, \quad \mu_1 = \mathbb{E}[G \mathbb{I}(G > 0)] = \frac{1}{\sqrt{2\pi}}, \quad \mu_*^2 := \mathbb{E}[\mathbb{I}(G > 0)] - \mu_0^2 - \mu_1^2 = \frac{1}{4} - \frac{1}{2\pi}.$$

Here, $\varphi(z)$ is the nonlinear component of the step function which is orthogonal to the constant and linear components in the following sense: $\mathbb{E}[\varphi(G)] = 0$ and $\mathbb{E}[G\varphi(G)] = 0$. We write

$$\mathbb{I}(\langle w_\ell, x_i \rangle > 0) = \frac{1}{2} + \frac{1}{\sqrt{2\pi}} \langle w_\ell, x_i \rangle + \mu_* u_{\ell i}, \quad \text{where: } u_{\ell i} := \varphi(\langle w_\ell, x_i \rangle),
 \tag{C.5}$$

noting that $\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle$ and $\langle \mathbf{w}_k, \mathbf{x}_i \rangle$ are jointly Gaussian with

$$\mathbb{E}(\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle^2) = \mathbb{E}(\langle \mathbf{w}_k, \mathbf{x}_i \rangle^2) = 1, \quad \mathbb{E}(\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle \langle \mathbf{w}_k, \mathbf{x}_i \rangle) = \langle \mathbf{w}_k, \mathbf{w}_\ell \rangle.$$

Therefore, we have (see e.g., [15, Table 1])

$$\begin{aligned} \mathbb{E}[\mathbb{I}(\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle > 0) \mathbb{I}(\langle \mathbf{w}_k, \mathbf{x}_i \rangle > 0)] &= \frac{\pi - \cos^{-1}(\langle \mathbf{w}_k, \mathbf{w}_\ell \rangle)}{2\pi} \\ &= \frac{1}{4} + \frac{1}{2\pi} \langle \mathbf{w}_\ell, \mathbf{w}_k \rangle + O(\langle \mathbf{w}_\ell, \mathbf{w}_k \rangle^3) \\ (C.6) \quad &= \frac{1}{4} + \frac{1}{2\pi} \langle \mathbf{w}_\ell, \mathbf{w}_k \rangle + O(d^{-3/2} \log^3(d)). \end{aligned}$$

To bound the correlation of variables $u_{\ell i}, u_{ki}$, we write

$$\begin{aligned} \mu_*^2 \mathbb{E}[u_{\ell i} u_{ki}] &= \mathbb{E} \left[\left\{ \mathbb{I}(\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle > 0) - \frac{1}{2} - \frac{1}{\sqrt{2\pi}} \langle \mathbf{w}_\ell, \mathbf{x}_i \rangle \right\} \left\{ \mathbb{I}(\langle \mathbf{w}_k, \mathbf{x}_i \rangle > 0) - \frac{1}{2} - \frac{1}{\sqrt{2\pi}} \langle \mathbf{w}_k, \mathbf{x}_i \rangle \right\} \right] \\ &= \mathbb{E}[\mathbb{I}(\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle > 0) \mathbb{I}(\langle \mathbf{w}_k, \mathbf{x}_i \rangle > 0)] + \mathbb{E} \left[\left(\frac{1}{2} + \frac{1}{\sqrt{2\pi}} \langle \mathbf{w}_\ell, \mathbf{x}_i \rangle \right) \left(\frac{1}{2} + \frac{1}{\sqrt{2\pi}} \langle \mathbf{w}_k, \mathbf{x}_i \rangle \right) \right] \\ (C.7) \quad &- \mathbb{E} \left[\mathbb{I}(\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle > 0) \left(\frac{1}{2} + \frac{1}{\sqrt{2\pi}} \langle \mathbf{w}_k, \mathbf{x}_i \rangle \right) \right] - \mathbb{E} \left[\mathbb{I}(\langle \mathbf{w}_k, \mathbf{x}_i \rangle > 0) \left(\frac{1}{2} + \frac{1}{\sqrt{2\pi}} \langle \mathbf{w}_\ell, \mathbf{x}_i \rangle \right) \right]. \end{aligned}$$

The first term above is calculated in (C.6). For the second term, we have

$$(C.8) \quad \mathbb{E} \left[\left(\frac{1}{2} + \frac{1}{\sqrt{2\pi}} \langle \mathbf{w}_\ell, \mathbf{x}_i \rangle \right) \left(\frac{1}{2} + \frac{1}{\sqrt{2\pi}} \langle \mathbf{w}_k, \mathbf{x}_i \rangle \right) \right] = \frac{1}{4} + \frac{\langle \mathbf{w}_\ell, \mathbf{w}_k \rangle}{2\pi}.$$

For the third term we write

$$\begin{aligned} \mathbb{E} \left[\mathbb{I}(\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle > 0) \left(\frac{1}{2} + \frac{1}{\sqrt{2\pi}} \langle \mathbf{w}_k, \mathbf{x}_i \rangle \right) \right] &= \frac{1}{4} + \frac{1}{\sqrt{2\pi}} \mathbb{E} \left[\mathbb{I}(\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle > 0) \langle \mathbf{w}_k, \mathbf{x}_i \rangle \right] \\ &= \frac{1}{4} + \frac{1}{\sqrt{2\pi}} \mathbb{E} \left[\langle \mathbf{w}_k, \mathbf{x}_i \rangle \mathbb{I}(\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle > 0) \right] \mathbb{P}(\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle > 0) \\ &\stackrel{(a)}{=} \frac{1}{4} + \frac{1}{2\sqrt{2\pi}} \langle \mathbf{w}_\ell, \mathbf{w}_k \rangle \frac{\phi(0)}{1 - \Phi(0)} \\ (C.9) \quad &= \frac{1}{4} + \frac{1}{2\pi} \langle \mathbf{w}_\ell, \mathbf{w}_k \rangle. \end{aligned}$$

Here (a) follows from lemma below.

Lemma C.3 For $Z_1, Z_2 \sim \mathcal{N}(0, 1)$ with $\mathbb{E}[Z_1 Z_2] = \rho$ we have

$$\mathbb{E}[Z_1 | Z_2 > z] = \rho \frac{\phi(z)}{(1 - \Phi(z))},$$

where $\phi(z) = \frac{1}{\sqrt{2\pi}} e^{-z^2/2}$ is the density of standard normal and $\Phi(z) = \int_{-\infty}^z \phi(t) dt$ is its CDF.

Using Equations (C.6), (C.8) and (C.9) in (C.7) we obtain

$$(C.10) \quad \mathbb{E}[u_{\ell i} u_{ki}] = O(d^{-3/2} \log^3(d)).$$

Substituting for the sign function $\mathbb{I}(\langle \mathbf{w}_\ell, \mathbf{x}_i \rangle > 0)$ from (C.5) we get

$$\begin{aligned} \mathbf{W}^\top \text{diag}(\mathbb{I}(\mathbf{W} \mathbf{x}_i > 0)) \boldsymbol{\theta} &= \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbb{I}(\mathbf{W} \mathbf{x}_i > 0) \\ (C.11) \quad &= \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \left(\frac{1}{2} \mathbf{1} + \mu_* \mathbf{u}_i + \frac{1}{\sqrt{2\pi}} \mathbf{W} \mathbf{x}_i \right), \end{aligned}$$

with $\mathbf{u}_i = (u_{\ell i})_{\ell=1}^N$. To lighten the notation, we use the shorthand $\mathbf{h}_i := \mathbf{W}^\top \text{diag}(\boldsymbol{\theta})(\frac{1}{2}\mathbf{1} + \mu_* \mathbf{u}_i)$. We next decompose $\eta_i(\boldsymbol{\theta})^2$ into three terms as follows:

$$(C.12) \quad \eta_i(\boldsymbol{\theta})^2 = \frac{1}{2\pi} \|\mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i\|_{\ell_2}^2 + \|\mathbf{h}_i\|_{\ell_2}^2 + \sqrt{\frac{2}{\pi}} \langle \mathbf{h}_i, \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i \rangle.$$

We next provide deviation bounds for each of these terms by putting which together we obtain the desired claim.

We start by the first term in (C.12). Note that since $\boldsymbol{\theta} \in \mathcal{C}_\theta$, we have the following bounds conditioned on the event $\mathcal{E}_\mathbf{W}$:

$$\begin{aligned} \|\mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W}\| &\leq \|\mathbf{W}\|^4 \|\boldsymbol{\theta}\|_{\ell_\infty}^2 = O(d^{-1} \log(d)), \\ \|\mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W}\|_F &\leq \sqrt{\min(d, N)} \|\mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W}\| = O(d^{-0.5}). \end{aligned}$$

Therefore, by applying Hanson-Wright's inequality [72] we get

$$(C.13) \quad \mathbb{P}\left\{\left|\|\mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i\|_{\ell_2}^2 - \mathbb{E}\left[\|\mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i\|_{\ell_2}^2\right]\right| > \gamma \log(d); \mathcal{E}_\mathbf{W}\right\} \leq 2e^{-c\gamma^2 d},$$

for an absolute constant $c > 0$.

For the second term we bound variation in the vector $\mathbf{h}_i$ itself from which we obtain a deviation bound on its norm $\|\mathbf{h}_i\|_{\ell_2}$. We write

$$(C.14) \quad \begin{aligned} \mathbb{E}\left[\|\mathbf{h}_i - \mathbb{E}[\mathbf{h}_i]\|_{\ell_2}^2\right] &= \mu_*^2 \mathbb{E}\left[\|\mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{u}_i\|_{\ell_2}^2\right] \\ &= \sum_{\ell, k} \langle \mathbf{w}_\ell, \mathbf{w}_k \rangle \theta_\ell \theta_k \mathbb{E}[u_{\ell i} u_{k i}] \leq \sum_{i, j} C \frac{1}{\sqrt{d}} \|\boldsymbol{\theta}\|_{\ell_\infty}^2 d^{-1.5} \log^3(d) = O(d^{-1} \log^4(d)), \end{aligned}$$

where we used the assumption $\boldsymbol{\theta} \in \mathcal{C}_\theta$ along with (C.10). We write $\mathbf{h}_i = \mathbb{E}[\mathbf{h}_i] + \boldsymbol{\delta}$ and define the event $\mathcal{E}_\delta := \{\|\boldsymbol{\delta}\|_{\ell_2} \leq \gamma\}$. Therefore by using Markov's inequality along with (C.14) we obtain $\mathbb{P}(\mathcal{E}_\delta) \geq 1 - c \frac{\log^4(d)}{d\gamma^2}$. Furthermore,

$$(C.15) \quad \|\mathbf{h}_i\|_{\ell_2}^2 = \|\mathbb{E}[\mathbf{h}_i]\|_{\ell_2}^2 + \|\boldsymbol{\delta}\|_{\ell_2}^2 + 2\langle \boldsymbol{\delta}, \mathbb{E}[\mathbf{h}_i] \rangle.$$

On the event $\mathcal{E}_\mathbf{W}$, we have $\|\mathbb{E}[\mathbf{h}_i]\|_{\ell_2} = \frac{1}{2} \|\mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{1}\|_{\ell_2} \leq \frac{1}{2} \|\mathbf{W}\| \|\boldsymbol{\theta}\|_\infty \sqrt{N} \leq C \sqrt{\log(d)}$, and so $|\langle \boldsymbol{\delta}, \mathbb{E}[\mathbf{h}_i] \rangle| \leq C \|\boldsymbol{\delta}\|_{\ell_2}$. Hence, on the event $\mathcal{E}_\mathbf{W} \cap \mathcal{E}_\delta$,

$$\left|\|\mathbf{h}_i\|_{\ell_2}^2 - \|\mathbb{E}[\mathbf{h}_i]\|_{\ell_2}^2\right| \leq \gamma^2 + 2C\sqrt{\log(d)}\gamma = O\left(\gamma\sqrt{\log(d)}\right).$$

This implies that $\|\mathbb{E}[\mathbf{h}_i]\|_{\ell_2}^2 = \mathbb{E}[\|\mathbf{h}_i\|_{\ell_2}^2] + O(\gamma \log(d))$, and therefore

$$(C.16) \quad \left|\|\mathbf{h}_i\|_{\ell_2}^2 - \mathbb{E}[\|\mathbf{h}_i\|_{\ell_2}^2]\right| = O\left(\gamma\sqrt{\log(d)}\right).$$

We next proceed to the third term in (C.12).

$$(C.17) \quad \begin{aligned} \langle \mathbf{h}_i, \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i \rangle &= \langle \mathbb{E}[\mathbf{h}_i], \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i \rangle + \langle \boldsymbol{\delta}, \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i \rangle \\ &= \frac{1}{2} \langle \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{1}, \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i \rangle + \langle \boldsymbol{\delta}, \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i \rangle. \end{aligned}$$

Note that on the event $\mathcal{E}_\mathbf{W}$ the first term above is a Lipschitz continuous function of the Gaussian vector $\mathbf{x}_i$ with Lipschitz constant

$$L = \|\mathbf{1}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W}\|_{\ell_2} \leq \sqrt{N} \|\boldsymbol{\theta}\|_{\ell_\infty}^2 \|\mathbf{W}\|^3 = O\left(\log(d)/\sqrt{d}\right).$$

By Gaussian isoperimetry [50], we have

(C.18)

$$\mathbb{P}\left(\left|\langle \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{1}, \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i \rangle - \mathbb{E}[\langle \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{1}, \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i \rangle]\right| \geq \gamma \log(d); \mathcal{E}_{\mathbf{W}}\right) \leq 2e^{-c\gamma^2 d},$$

for some constant $c > 0$. For the second term of (C.17), note that on the event $\mathcal{E}_{\delta} \cap \mathcal{E}_{\mathbf{W}} \cap \mathcal{E}_{\mathbf{x}}$,

(C.19)

$$|\langle \boldsymbol{\delta}, \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i \rangle| \leq \|\boldsymbol{\delta}\|_{\ell_2} \|\mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i\|_{\ell_2} \leq \|\boldsymbol{\delta}\|_{\ell_2} \|\mathbf{W}\|^2 \|\boldsymbol{\theta}\|_{\ell_\infty} \|\mathbf{x}_i\|_{\ell_2} = O(\gamma \log(d)).$$

Combining (C.18) and (C.19) with the decomposition (C.17) we get that on the event $\mathcal{E}_{\delta} \cap \mathcal{E}_{\mathbf{W}} \cap \mathcal{E}_{\mathbf{x}}$,

$$(C.20) \quad \left| \langle \mathbf{h}_i, \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i \rangle - \mathbb{E}[\langle \mathbf{h}_i, \mathbf{W}^\top \text{diag}(\boldsymbol{\theta}) \mathbf{W} \mathbf{x}_i \rangle] \right| = O(\gamma \log(d)),$$

with probability at least $1 - 2e^{-c\gamma^2 d}$. Putting together the deviation bounds for the three terms, given by (C.13), (C.16) and (C.20), we arrive at

$$\mathbb{P}\left(|\eta_i(\boldsymbol{\theta})^2 - \mathbb{E}[\eta_i(\boldsymbol{\theta})^2]| > C\gamma \log(d); \mathcal{E}_{\mathbf{W}} \cap \mathcal{E}_{\mathbf{x}}\right) \leq 4e^{-c\gamma^2 d} + \frac{c \log^4(d)}{d\gamma^2} = O\left(\frac{\log^4(d)}{d\gamma^2}\right).$$

Note that the above relation holds for any $\gamma > 0$. The result of the lemma now follows by letting $\gamma \leftarrow C\gamma \log(d)$.

C.3. Proof of Lemma C.2

Define the matrix

$$\mathbf{A} := \text{diag}(\mathbb{1}(\mathbf{W} \mathbf{x}_i > 0)) \mathbf{W} \mathbf{W}^\top \text{diag}(\mathbb{1}(\mathbf{W} \mathbf{x}_i > 0)).$$

By triangle inequality we have

$$\begin{aligned} |\eta_i(\boldsymbol{\theta})^2 - \eta_i(\tilde{\boldsymbol{\theta}})^2| &= |\langle \mathbf{A} \boldsymbol{\theta}, \boldsymbol{\theta} \rangle - \langle \mathbf{A} \tilde{\boldsymbol{\theta}}, \tilde{\boldsymbol{\theta}} \rangle| \\ &= |\langle \mathbf{A} \boldsymbol{\theta}, \boldsymbol{\theta} - \tilde{\boldsymbol{\theta}} \rangle + \langle \mathbf{A}(\boldsymbol{\theta} - \tilde{\boldsymbol{\theta}}), \tilde{\boldsymbol{\theta}} \rangle| \\ &\leq \|\mathbf{A}\| \|\boldsymbol{\theta}\|_{\ell_2} \|\boldsymbol{\theta} - \tilde{\boldsymbol{\theta}}\|_{\ell_2} + \|\mathbf{A}\| \|\boldsymbol{\theta} - \tilde{\boldsymbol{\theta}}\|_{\ell_2} \|\tilde{\boldsymbol{\theta}}\|_{\ell_2} \leq 2c_0\gamma \|\mathbf{A}\|, \end{aligned}$$

where in the last step we used the fact that $\boldsymbol{\theta}, \tilde{\boldsymbol{\theta}} \in \mathcal{S}^{d-1}$ and $\|\boldsymbol{\theta} - \tilde{\boldsymbol{\theta}}\|_{\ell_2} \leq c_0\gamma$. So it suffices to bound $\|\mathbf{A}\|$. By definition, the matrix $\mathbf{A}$ is obtained by selecting a subset of rows and columns of $\mathbf{W} \mathbf{W}^\top$ and replacing them with zeros. Therefore, $\|\mathbf{A}\| \leq \|\mathbf{W} \mathbf{W}^\top\| \leq \sqrt{\psi_{1,d}} + C$, on the event $\mathcal{E}_{\mathbf{W}}$.

Since the above bound in Lemma C.2 holds for any vector $\mathbf{x}_i$, a similar bound also holds if the terms are replaced by their expectation with respect to $\mathbf{x}_i$, whence we obtain $|\mathbb{E}[\eta_i(\boldsymbol{\theta})^2] - \mathbb{E}[\eta_i(\tilde{\boldsymbol{\theta}})^2]| \leq 2c_0\gamma(\sqrt{\psi_{1,d}} + C)$.

By definition of $\nu_i(\boldsymbol{\theta}; \gamma)$ we have

$$\begin{aligned} \nu_i(\boldsymbol{\theta}; \gamma) = 1 &\implies |\eta_i(\boldsymbol{\theta})^2 - \mathbb{E}[\eta_i(\boldsymbol{\theta})^2]| \geq \gamma \\ &\implies |\eta_i(\tilde{\boldsymbol{\theta}})^2 - \mathbb{E}[\eta_i(\tilde{\boldsymbol{\theta}})^2]| + |(\eta_i(\boldsymbol{\theta})^2 - \mathbb{E}[\eta_i(\boldsymbol{\theta})^2]) - (\eta_i(\tilde{\boldsymbol{\theta}})^2 - \mathbb{E}[\eta_i(\tilde{\boldsymbol{\theta}})^2])| \geq \gamma \\ &\implies |\eta_i(\tilde{\boldsymbol{\theta}})^2 - \mathbb{E}[\eta_i(\tilde{\boldsymbol{\theta}})^2]| \geq \gamma - 2c_0\gamma(\sqrt{\psi_{1,d}} + C) \\ (C.21) \quad &\implies \nu_i(\tilde{\boldsymbol{\theta}}; \gamma(1 - 2c_0\sqrt{\psi_{1,d}} - 2c_0C)) = 1. \end{aligned}$$

C.3.1. *Proof of Lemma C.3* The conditional distribution of Z_1 given Z_2 is

$$Z_1|Z_2 = z_2 \sim \mathcal{N}(\rho z_2, (1 - \rho^2)).$$

Therefore, $\mathbb{E}[Z_1|Z_2 = z_2] = \rho z_2$ and so

$$\mathbb{E}[Z_1|Z_2 > z] = \mathbb{E}[Z_1|Z_2 = z_2] \mathbb{P}(Z_2 = z_2|Z_2 > z) dz_2 = \rho \mathbb{E}[Z_2|Z_2 > z].$$

Using the properties of the expectation of a truncated normal distribution, we have

$$\mathbb{E}[Z_2|Z_2 > z] = \frac{\phi(z)}{(1 - \Phi(z))},$$

which completes the proof.

C.4. Proof of Lemma 6.6 By (6.6) it suffices to show that

$$\sup_{\theta \in \mathcal{C}_\theta} \frac{|\mathring{\mathcal{L}}(\theta) - \mathring{\mathcal{L}}(\theta)|}{1 + \min(\mathring{\mathcal{L}}(\theta), \mathring{\mathcal{L}}(\theta))} = o_{d,\mathbb{P}}(1).$$

To lighten the notation we define the shorthand $\alpha_i := |y_i - \theta^\top \sigma(\mathbf{W} \mathbf{x}_i)| + \varepsilon \|\mathbf{W}^\top \text{diag}(\mathbf{W} \mathbf{x}_i > 0) \theta\|_{\ell_2}$ and so $\mathring{\mathcal{L}}(\theta) = 1/(2n) \sum_{i=1}^n \alpha_i^2 + \frac{\zeta}{2} \theta^\top \Omega \theta$. We write $\mathring{\mathcal{L}}(\theta) = 1/(2n) \sum_{i=1}^n (\alpha_i + \beta_i)^2 + \frac{\zeta}{2} \theta^\top \Omega \theta$ with

$$\beta_i := \varepsilon \|\mathbf{J} \theta\|_{\ell_2} - \varepsilon \|\mathbf{W}^\top \text{diag}(\mathbf{W} \mathbf{x}_i > 0) \theta\|_{\ell_2}.$$

Since for any two positive values a, b we have $|a - b| \leq \sqrt{|a^2 - b^2|}$, we can write

$$(C.22) \quad |\beta_i| \leq \varepsilon \left(\|\mathbf{W}^\top \text{diag}(\mathbf{W} \mathbf{x}_i > 0) \theta\|_{\ell_2}^2 - \|\mathbf{J} \theta\|_{\ell_2}^2 \right)^{1/2} = [\eta_i(\theta)^2 - \mathbb{E}[\eta_i(\theta)^2]]^{1/2}.$$

Note that on the event $\mathcal{E}_{\mathbf{W}}$, $\|\mathbf{W}\|$ is bounded and so $\|\mathbf{W}^\top \text{diag}(\mathbf{W} \mathbf{x}_i > 0)\|$ is also bounded. Since $\theta \in \mathcal{C}_\theta$, we have $\|\theta\|_{\ell_2} = O(1)$, which along with Lemma F.3 imply that $\max_{i \in [n]} |\eta_i(\theta)|$ and $\max_{i \in [n]} |\mathbb{E}[\eta_i(\theta)]|$ are both $O_{d,\mathbb{P}}(1)$. Therefore: (i) defining $\beta = (\beta_1, \dots, \beta_n)$, we have $\|\beta\|_{\ell_2} = O_{d,\mathbb{P}}(1)$; (ii) Using (C.22) along with Corollary 6.5, we get $\frac{1}{n} |\{i : |\beta_i| > \frac{1}{\sqrt{\log(d)}}\}| = o_{d,\mathbb{P}}(1)$.

From the above we can deduce that $\frac{1}{n} \|\beta\|_{\ell_2}^2 = o_{d,\mathbb{P}}(1)$. We also define $\alpha = (\alpha_1, \dots, \alpha_n)$. For any $\theta \in \mathcal{C}_\theta$ we have

$$\begin{aligned} |\mathring{\mathcal{L}}(\theta) - \mathring{\mathcal{L}}(\theta)| &= \left| \frac{1}{2n} \sum_{i=1}^n (\alpha_i + \beta_i)^2 - \sum_{i=1}^n \alpha_i^2 \right| \\ &= \frac{\|\beta\|_{\ell_2}^2}{2n} + \frac{1}{n} \left| \sum_{i=1}^n \alpha_i \beta_i \right| \\ &\leq \frac{\|\beta\|_{\ell_2}^2}{2n} + \frac{1}{n} \|\beta\|_{\ell_2} \|\alpha\|_{\ell_2} \\ &\leq \frac{\|\beta\|_{\ell_2}^2}{2n} + \frac{\|\beta\|_{\ell_2}}{\sqrt{n}} \frac{\|\alpha\|_{\ell_2}}{\sqrt{n}} \\ &\leq \frac{\|\beta\|_{\ell_2}^2}{2n} + \frac{\|\beta\|_{\ell_2}}{2\sqrt{n}} \left(1 + \frac{\|\alpha\|_{\ell_2}^2}{n} \right) \end{aligned}$$

$$\begin{aligned}
&\leq \frac{\|\beta\|_{\ell_2}^2}{2n} + \frac{\|\beta\|_{\ell_2}}{2\sqrt{n}} + \frac{\|\beta\|_{\ell_2}}{\sqrt{n}} \mathring{\mathcal{L}}(\theta) \\
&= o_{d,\mathbb{P}}(1)(1 + \mathring{\mathcal{L}}(\theta)),
\end{aligned}$$

where the last step holds because $\frac{\|\beta\|_{\ell_2}}{\sqrt{n}} = o_{d,\mathbb{P}}(1)$.

By a similar argument, we also get

$$|\mathring{\mathring{\mathcal{L}}}(\theta) - \mathring{\mathcal{L}}(\theta)| \leq o_{d,\mathbb{P}}(1)(1 + \mathring{\mathring{\mathcal{L}}}(\theta)).$$

Combining these two bounds we get $|\mathring{\mathring{\mathcal{L}}}(\theta) - \mathring{\mathcal{L}}(\theta)| \leq o_{d,\mathbb{P}}(1) \left(1 + \min(\mathring{\mathcal{L}}(\theta), \mathring{\mathring{\mathcal{L}}}(\theta))\right)$.

We next proceed to the second part. By optimality of $\widehat{\theta}$ and $\mathring{\mathring{\theta}}$ we have

$$(C.23) \quad \mathring{\mathring{\mathcal{L}}}(\widehat{\theta}) < \mathring{\mathring{\mathcal{L}}}(\mathbf{0}) = \frac{1}{n} \sum_{i=1}^n y_i^2 = O_{d,\mathbb{P}}(1), \quad \mathcal{L}(\widehat{\theta}) < \mathcal{L}(\mathbf{0}) = \frac{1}{n} \sum_{i=1}^n y_i^2 = O_{d,\mathbb{P}}(1).$$

As shown in Proposition 6.2, $\widehat{\theta} \in \mathcal{C}_\theta$ with high probability. Likewise we have $\widehat{\theta}^* \in \mathcal{C}_\theta$, with high probability (this follows from Lemma D.7 for the special case of $k = n$ in that lemma.)

Therefore using the first part of the current lemma,

$$|\mathring{\mathring{\mathcal{L}}}(\widehat{\theta}) - \mathcal{L}(\widehat{\theta})| = o_{d,\mathbb{P}}(1), \quad |\mathring{\mathring{\mathcal{L}}}(\widehat{\theta}^*) - \mathcal{L}(\widehat{\theta}^*)| = o_{d,\mathbb{P}}(1).$$

We therefore obtain

$$0 \leq \mathcal{L}(\widehat{\theta}^*) - \mathcal{L}(\widehat{\theta}) < (\mathcal{L}(\widehat{\theta}^*) - \mathring{\mathring{\mathcal{L}}}(\widehat{\theta}^*)) + \underbrace{(\mathring{\mathring{\mathcal{L}}}(\widehat{\theta}^*) - \mathring{\mathring{\mathcal{L}}}(\widehat{\theta}))}_{\leq 0} + (\mathring{\mathring{\mathcal{L}}}(\widehat{\theta}) - \mathcal{L}(\widehat{\theta})) \leq o_{d,\mathbb{P}}(1).$$

Since $\mathcal{L}(\theta)$ is $\frac{\zeta}{2}$ -strongly convex we have

$$\|\widehat{\theta}^* - \widehat{\theta}\|_{\ell_2} \leq o_{d,\mathbb{P}}(1)/\zeta \rightarrow 0, \text{ as } d \rightarrow \infty.$$

C.5. Proof of Lemma 6.7 We define

$$\mathring{\mathring{\mathcal{A}}}R(\theta) := \mathbb{E} \left[\left(|y - \theta^\top \sigma(\mathbf{W}\mathbf{x})| + \varepsilon_{\text{test}} \|\mathbf{W}^\top \text{diag} \mathbb{I}(\mathbf{W}\mathbf{x} > 0) \theta\|_{\ell_2} \right)^2 \right].$$

As an immediate result of Proposition 6.1, we have $\sup_{\theta \in \mathcal{C}_\theta} |\mathring{\mathring{\mathcal{A}}}R(\theta) - \mathring{\mathcal{A}}R(\theta)| = o_d(1)$. Therefore, it suffices to show that

$$(C.24) \quad \sup_{\theta \in \mathcal{C}_\theta} \frac{|\mathring{\mathring{\mathcal{A}}}R(\theta) - \mathring{\mathcal{A}}R(\theta)|}{\sqrt{\mathring{\mathcal{A}}R(\theta)}} = o_{d,\mathbb{P}}(1).$$

By expanding the terms in $\mathring{\mathring{\mathcal{A}}}R(\theta)$ and invoking our notation $\eta_i(\theta) = \|\mathbf{W}^\top \text{diag} \mathbb{I}(\mathbf{W}\mathbf{x} > 0) \theta\|_{\ell_2}$, we have

$$\mathring{\mathring{\mathcal{A}}}R(\theta) = \mathbb{E}[(y - \theta^\top \sigma(\mathbf{W}\mathbf{x}))^2] + \varepsilon_{\text{test}}^2 \mathbb{E}[\eta_i(\theta)^2] + 2\varepsilon_{\text{test}} \mathbb{E}[|y - \theta^\top \sigma(\mathbf{W}\mathbf{x})| \eta_i(\theta)].$$

Likewise we have

$$\mathring{\mathcal{A}}R(\theta) = \mathbb{E}[(y - \theta^\top \sigma(\mathbf{W}\mathbf{x}))^2] + \varepsilon_{\text{test}}^2 \|\mathbf{J}\theta\|_{\ell_2}^2 + 2\varepsilon_{\text{test}} \mathbb{E}[|y - \theta^\top \sigma(\mathbf{W}\mathbf{x})|] \|\mathbf{J}\theta\|_{\ell_2}.$$

Recall that by definition of $\mathbf{J}$ we have $\mathbb{E}[\eta_i(\boldsymbol{\theta})^2] = \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2$. Hence,

$$\begin{aligned}
 |\mathring{\text{AR}}(\boldsymbol{\theta}) - \mathring{\mathring{\text{AR}}}(\boldsymbol{\theta})| &= 2\varepsilon_{\text{test}} \left| \mathbb{E} \left[|y - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x})| (\eta_i(\boldsymbol{\theta}) - \sqrt{\mathbb{E}[\eta_i(\boldsymbol{\theta})^2]}) \right] \right| \\
 &\leq 2\varepsilon_{\text{test}} \mathbb{E} \left[(y - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}))^2 \right]^{1/2} \mathbb{E} \left[(\eta_i(\boldsymbol{\theta}) - \sqrt{\mathbb{E}[\eta_i(\boldsymbol{\theta})^2]})^2 \right]^{1/2} \\
 (C.25) \quad &\leq 2\varepsilon_{\text{test}} \sqrt{\mathring{\text{AR}}(\boldsymbol{\theta})} \mathbb{E} \left[(\eta_i(\boldsymbol{\theta}) - \sqrt{\mathbb{E}[\eta_i(\boldsymbol{\theta})^2]})^2 \right]^{1/2}.
 \end{aligned}$$

To bound the right-hand side, note that for any two positive values a, b we have $(a - b)^2 \leq |a^2 - b^2|$. Therefore,

$$(C.26) \quad \mathbb{E} \left[(\eta_i(\boldsymbol{\theta}) - \sqrt{\mathbb{E}[\eta_i(\boldsymbol{\theta})^2]})^2 \right] \leq \mathbb{E} \left[|\eta_i(\boldsymbol{\theta})^2 - \mathbb{E}[\eta_i(\boldsymbol{\theta})^2]| \right].$$

Also recall that for any non-negative random variable Z , we have $\mathbb{E}[Z] = \int_0^\infty \mathbb{P}(Z \geq z) dz$. Therefore,

$$\begin{aligned}
 \mathbb{E} \left[|\eta_i(\boldsymbol{\theta})^2 - \mathbb{E}[\eta_i(\boldsymbol{\theta})^2]| \right] &= \mathbb{E} \left[|\eta_i(\boldsymbol{\theta})^2 - \mathbb{E}[\eta_i(\boldsymbol{\theta})^2]|; \mathcal{E}_{\mathbf{W}} \cap \mathcal{E}_{\mathbf{x}} \right] + \mathbb{P}((\mathcal{E}_{\mathbf{W}} \cap \mathcal{E}_{\mathbf{x}})^c) \\
 &\quad \int_0^\infty \mathbb{P} \left(|\eta_i(\boldsymbol{\theta})^2 - \mathbb{E}[\eta_i(\boldsymbol{\theta})^2]| \geq \gamma \right) d\gamma + \mathbb{P}((\mathcal{E}_{\mathbf{W}} \cap \mathcal{E}_{\mathbf{x}})^c) \\
 (C.27) \quad &\leq \int_0^\infty \min \left(\frac{c}{d\gamma^2}, 1 \right) d\gamma + c \exp(-\log^2(d)/c) + ne^{-d}
 \end{aligned}$$

where the inequality follows from Lemma C.1. Next, we have

$$\begin{aligned}
 \int_0^\infty \min \left(\frac{c \log^6(d)}{d\gamma^2}, 1 \right) d\gamma &= \int_0^{\sqrt{c \log^6(d)/d}} d\gamma + \int_{\sqrt{c \log^6(d)/d}}^\infty \frac{c \log^6(d)}{d\gamma^2} d\gamma \\
 (C.28) \quad &= \sqrt{\frac{c \log^6(d)}{d}} + \frac{c \log^6(d)}{d} \sqrt{\frac{d}{c \log^6(d)}} = 2\sqrt{\frac{c \log^6(d)}{d}}.
 \end{aligned}$$

Combining Eqs. (C.26), (C.27) and (C.28) we arrive at

$$\mathbb{E} \left[(\eta_i(\boldsymbol{\theta}) - \sqrt{\mathbb{E}[\eta_i(\boldsymbol{\theta})^2]})^2 \right] \leq 2\sqrt{\frac{c \log^6(d)}{d}} + c \exp(-\log^2(d)/c) + ne^{-d} = o_d(\log^3(d)d^{-1/2}).$$

Using the above bound in (C.25) we get that uniformly over $\boldsymbol{\theta} \in \mathcal{C}_{\boldsymbol{\theta}}$,

$$|\mathring{\text{AR}}(\boldsymbol{\theta}) - \mathring{\mathring{\text{AR}}}(\boldsymbol{\theta})| \leq \sqrt{\mathring{\text{AR}}(\boldsymbol{\theta})} o_{d,\mathbb{P}}(1).$$

This completes the proof of claim (C.24).

APPENDIX D: PROOFS OF STEP 3: THE GAUSSIAN EQUIVALENCE PROPERTY

D.1. Proof of Proposition 6.8 As proved in [30, Theorem 2], under the assumptions of Proposition 6.8, $(\boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}), \beta^\top \mathbf{x}_0)$ converges in distribution to $(\boldsymbol{\theta}^\top \mathbf{f}, \beta^\top \mathbf{x})$. We first show that $\mathring{\mathring{\text{AR}}}(\boldsymbol{\theta}) = \mathring{\mathring{\text{AR}}}_{\text{nl}}(\boldsymbol{\theta}) + o_d(1)$. Recalling the definition of $\mathring{\mathring{\text{AR}}}_{\text{nl}}(\boldsymbol{\theta})$ given by (6.18), and plugging for $y = \beta^\top \mathbf{x} + \xi$, we write

$$\begin{aligned}
 \mathring{\mathring{\text{AR}}}_{\text{nl}}(\boldsymbol{\theta}) &= \mathbb{E} \left[(|\beta^\top \mathbf{x} - \boldsymbol{\theta}^\top \mathbf{f} + \xi| + \varepsilon \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2})^2 \right] \\
 (D.1) \quad &= \mathbb{E} \left[(\beta^\top \mathbf{x} - \boldsymbol{\theta}^\top \mathbf{f} + \xi)^2 \right] + \varepsilon \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2} \mathbb{E} [|\beta^\top \mathbf{x} - \boldsymbol{\theta}^\top \mathbf{f} + \xi|] + \varepsilon^2 \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2
 \end{aligned}$$

Therefore, $\overset{\circ\circ}{\text{AR}}_{\text{nl}}(\boldsymbol{\theta})$ can be written in terms of the first and second moment of random variable $|\boldsymbol{\beta}^\top \mathbf{x} - \boldsymbol{\theta}^\top \mathbf{f} + \xi|$ which converges in distribution to $|\boldsymbol{\beta}^\top \mathbf{x} - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}) + \xi|$. To show that $\overset{\circ\circ}{\text{AR}}_{\text{nl}}(\boldsymbol{\theta}) - \overset{\circ\circ}{\text{AR}}(\boldsymbol{\theta}) \rightarrow 0$ as $d \rightarrow \infty$, we need to show that the first and second moments of $|\boldsymbol{\beta}^\top \mathbf{x} - \boldsymbol{\theta}^\top \mathbf{f} + \xi|$ converge respectively to the first and second moments of $|\boldsymbol{\beta}^\top \mathbf{x} - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}) + \xi|$. As an application of [7, Corollary of Theorem 25.12], it suffices to show that $|\boldsymbol{\beta}^\top \mathbf{x} - \boldsymbol{\theta}^\top \mathbf{f} + \xi|$ has bounded third moment. To show this, note that by Holder's inequality, $|a + b + c|^3 \leq 3(|a|^3 + |b|^3 + |c|^3)$. Furthermore, $\mathbb{E}[|\xi|^3] = 2$, $\mathbb{E}[|\boldsymbol{\beta}^\top \mathbf{x}|^3] = 2 \|\boldsymbol{\beta}\|_{\ell_2}^3 = 2$. Hence,

$$(D.2) \quad \mathbb{E}\left[|\boldsymbol{\beta}^\top \mathbf{x} - \boldsymbol{\theta}^\top \mathbf{f} + \xi|^3\right] \leq 3(16 + \mathbb{E}[|\boldsymbol{\theta}^\top \mathbf{f}|^3]).$$

By using the Holder's inequality again we have

$$(D.3) \quad \begin{aligned} \mathbb{E}[|\boldsymbol{\theta}^\top \mathbf{f}|^3] &= \mathbb{E}\left[\left|\frac{1}{\sqrt{2\pi}} \mathbf{1}^\top \boldsymbol{\theta} + \frac{1}{2} \boldsymbol{\theta}^\top \mathbf{W}\mathbf{x} + \sqrt{\frac{1}{4} - \frac{1}{2\pi}} \boldsymbol{\theta}^\top \mathbf{u}\right|^3\right] \\ &\leq 3 \left(\frac{1}{\sqrt{2\pi}^3} (\mathbf{1}^\top \boldsymbol{\theta})^3 + \frac{1}{4} \|\mathbf{W}^\top \boldsymbol{\theta}\|_{\ell_2}^3 + 2 \left(\frac{1}{4} - \frac{1}{2\pi}\right)^{3/2} \|\boldsymbol{\theta}\|_{\ell_2}^3 \right). \end{aligned}$$

Now note that by our assumption $\overset{\circ\circ}{\text{AR}}_{\text{nl}}(\boldsymbol{\theta})$ is bounded, which in conjunction with characterization (6.20) implies that $\|\boldsymbol{\theta}\|_{\ell_2}$, $\mathbf{1}^\top \boldsymbol{\theta}$ and $\|\mathbf{W}^\top \boldsymbol{\theta} - 2\boldsymbol{\beta}\|_{\ell_2}^2$ are bounded as $d \rightarrow \infty$. This also implies that $\|\mathbf{W}^\top \boldsymbol{\theta}\|_{\ell_2}^2 \leq (\|\mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta}\|_{\ell_2} + \|\boldsymbol{\beta}\|_{\ell_2})^2 \leq (\|\mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta}\|_{\ell_2} + 1)^2$ is bounded. Putting these together, we obtain that $\mathbb{E}\left[|\boldsymbol{\beta}^\top \mathbf{x} - \boldsymbol{\theta}^\top \mathbf{f} + \xi|^3\right]$ is bounded, which completes the argument for showing that $\overset{\circ\circ}{\text{AR}}_{\text{nl}}(\boldsymbol{\theta}) = \overset{\circ\circ}{\text{AR}}(\boldsymbol{\theta}) + o_d(1)$.

We next prove the characterization (6.20). Note that

$$(D.4) \quad \begin{aligned} y - \boldsymbol{\theta}^\top \mathbf{f} &= \xi + \boldsymbol{\beta}^\top \mathbf{x} - \frac{1}{2} \boldsymbol{\theta}^\top \mathbf{W}\mathbf{x} - \sqrt{\frac{1}{4} - \frac{1}{2\pi}} \boldsymbol{\theta}^\top \mathbf{u} \\ &= \xi + \langle \boldsymbol{\beta} - \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta}, \mathbf{x} \rangle - \sqrt{\frac{1}{4} - \frac{1}{2\pi}} \boldsymbol{\theta}^\top \mathbf{u} \\ &\sim \text{N}(0, M(\boldsymbol{\theta})^2), \end{aligned}$$

with

$$M(\boldsymbol{\theta})^2 = \tau^2 + \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \left(\frac{1}{4} - \frac{1}{2\pi} \right) \|\boldsymbol{\theta}\|_{\ell_2}^2.$$

We then write

$$(D.5) \quad \begin{aligned} \overset{\circ\circ}{\text{AR}}_{\text{nl}}(\boldsymbol{\theta}) &= \mathbb{E}\left[\left(|y - \boldsymbol{\theta}^\top \mathbf{f}| + \varepsilon_{\text{test}} \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}\right)^2\right] \\ &= \mathbb{E}\left[\left(y - \boldsymbol{\theta}^\top \mathbf{f}\right)^2\right] + 2\varepsilon_{\text{test}} \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2} \mathbb{E}\left[|y - \boldsymbol{\theta}^\top \mathbf{f}|\right] + \varepsilon_{\text{test}}^2 \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2 \\ &= M(\boldsymbol{\theta})^2 + \varepsilon_{\text{test}}^2 \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2 + 2\sqrt{\frac{2}{\pi}} \varepsilon_{\text{test}} M(\boldsymbol{\theta}) \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}, \end{aligned}$$

using the first and second moment of folded normal distribution.

D.2. Proof of Theorem 6.9 Before stating the proof, we remark that our proof is an adaptation of the powerful machinery developed [39]; however, since our individual adversarial losses have an additional term $\varepsilon \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}$, there are some additional details in the proofs which we provide in the following. Also, since our activation function is not odd, we will use the CLT-type result of [30] instead of the one provided in [39].

Recall that we are seeking to analyze the asymptotic values of the following quantities:

$$\Phi_A := \min_{\boldsymbol{\theta}} \frac{1}{n} \sum_{i=1}^n \left(|y_i - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}_i)| + \varepsilon \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2} \right)^2 + \lambda \|\boldsymbol{\theta}\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \lambda_s \frac{\log(d)}{d} (\mathbf{1}^\top \boldsymbol{\theta})^2,$$

$$\Phi_B := \min_{\boldsymbol{\theta}} \frac{1}{n} \sum_{i=1}^n \left(|y_i - \boldsymbol{\theta}^\top \mathbf{f}_i| + \varepsilon \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2} \right)^2 + \lambda \|\boldsymbol{\theta}\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \lambda_s \frac{\log(d)}{d} (\mathbf{1}^\top \boldsymbol{\theta})^2,$$

To simplify our notation, and without loss of generality, we absorb the value ε into $\mathbf{J}$ and, with a slight abuse of notation, consider $\mathbf{J} \leftarrow \varepsilon \mathbf{J}$ (and hence the eigenvalues of the matrix $\mathbf{J}$ depend on ε). Hence, above the quantities of interest become

$$\Phi_A := \min_{\boldsymbol{\theta}} \frac{1}{n} \sum_{i=1}^n \left(|y_i - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}_i)| + \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2} \right)^2 + \lambda \|\boldsymbol{\theta}\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \lambda_s \frac{\log(d)}{d} (\mathbf{1}^\top \boldsymbol{\theta})^2,$$

$$\Phi_B := \min_{\boldsymbol{\theta}} \frac{1}{n} \sum_{i=1}^n \left(|y_i - \boldsymbol{\theta}^\top \mathbf{f}_i| + \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2} \right)^2 + \lambda \|\boldsymbol{\theta}\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \lambda_s \frac{\log(d)}{d} (\mathbf{1}^\top \boldsymbol{\theta})^2,$$

For technical reasons, we first need to make the objectives smooth. We thus define $g(x) = \sqrt{x + \gamma}$, and define the smoothed loss

$$(D.6) \quad \ell(\boldsymbol{\theta}; \mathbf{r}, y) = (\boldsymbol{\theta}^\top \mathbf{r} - y)^2 + \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2 + 2g\left((\boldsymbol{\theta}^\top \mathbf{r} - y)^2 \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2\right),$$

Note that when $\gamma = 0$, we have $\ell(\boldsymbol{\theta}; \sigma(\mathbf{W}\mathbf{x}_i), y_i) = (|y_i - \boldsymbol{\theta}^\top \sigma(\mathbf{W}\mathbf{x}_i)| + \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2})^2$ and $\ell(\boldsymbol{\theta}; \mathbf{f}_i, y_i) = (|y_i - \boldsymbol{\theta}^\top \mathbf{f}_i| + \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2})^2$.

In the following, we consider an arbitrary but *fixed* value $\gamma > 0$, and with some abuse of notation, we let

$$\Phi_A := \min_{\boldsymbol{\theta}} \frac{1}{n} \sum_{i=1}^n \ell(\boldsymbol{\theta}; \sigma(\mathbf{W}\mathbf{x}_i), y_i) + \lambda \|\boldsymbol{\theta}\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \lambda_s \frac{\log(d)}{d} (\mathbf{1}^\top \boldsymbol{\theta})^2,$$

$$\Phi_B := \min_{\boldsymbol{\theta}} \frac{1}{n} \sum_{i=1}^n \ell(\boldsymbol{\theta}; \mathbf{f}_i, y_i) + \lambda \|\boldsymbol{\theta}\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \lambda_s \frac{\log(d)}{d} (\mathbf{1}^\top \boldsymbol{\theta})^2,$$

For these quantities Φ_A, Φ_B , we show then show in the following that the statement of the Theorem is true. Then, the result of the Theorem for the original losses – i.e. when $\gamma = 0$ – follows simply by taking the limit $\gamma \rightarrow 0$ (note that this limit is taken *after* the limit $d \rightarrow \infty$; also, note that $\sup_{x \geq 0} \{g(x) - x\} = \sqrt{\gamma}$).

We will use the Lindeberg's leave-one-out technique. In a nutshell, we start with the quantity Φ_B , and through n consecutive steps, we reach to the quantity Φ_A . In the k -th step, we will replace the feature vector $\mathbf{f}_k$ with $\sigma(\mathbf{W}\mathbf{x}_k)$. We will then show that each of these replacements has a negligible effect (i.e. $o_n(1)/n$) on our quantities of interest, leading to the proof of the theorem.

Let us now proceed with the details. The proof has multiple steps which will be put together in Section D.2.6 to obtain the proof of the theorem.

We begin by defining

$$(D.7) \quad \Phi_k = \min_{\boldsymbol{\theta}} \frac{1}{n} \sum_{i=1}^k \ell(\boldsymbol{\theta}; \sigma(\mathbf{W}\mathbf{x}_i), y_i) + \frac{1}{n} \sum_{i=k+1}^n \ell(\boldsymbol{\theta}; \mathbf{f}_i, y_i) + \lambda \|\boldsymbol{\theta}\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \lambda_s \frac{\log(d)}{d} (\mathbf{1}^\top \boldsymbol{\theta})^2,$$

Roughly speaking, our goal is to show that for all $k \in [n]$, we have $\Phi_k \approx \Phi_{k-1} + o_n(1)/n$. To make this entirely rigorous, we need to define several new quantities and understand their relations. For $k \in [n]$ let

(D.8)

$$R_{-k}(\boldsymbol{\theta}) = \frac{1}{n} \sum_{i=1}^{k-1} \ell(\boldsymbol{\theta}; \sigma(\mathbf{W} \mathbf{x}_i), y_i) + \frac{1}{n} \sum_{i=k+1}^n \ell(\boldsymbol{\theta}; \mathbf{f}_i, y_i) + \lambda \|\boldsymbol{\theta}\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \lambda_s \frac{\log(d)}{d} (\mathbf{1}^\top \boldsymbol{\theta})^2,$$

and

$$(D.9) \quad R_k(\boldsymbol{\theta}, \mathbf{r}) = \frac{1}{n} \ell(\boldsymbol{\theta}; \mathbf{r}, y_k) + R_{-k}(\boldsymbol{\theta}).$$

Let us denote the minimizers of the above two objectives by

$$(D.10) \quad \boldsymbol{\theta}_{-k}^* = \arg \min_{\boldsymbol{\theta}} R_{-k}(\boldsymbol{\theta}), \text{ and } \boldsymbol{\theta}_k^*(\mathbf{r}) = \arg \min_{\boldsymbol{\theta}} R_k(\boldsymbol{\theta}, \mathbf{r})$$

and

$$(D.11) \quad \Phi_{-k} = \min_{\boldsymbol{\theta}} R_{-k}(\boldsymbol{\theta}),$$

and

$$(D.12) \quad \Phi_k(\mathbf{r}) = \min_{\boldsymbol{\theta}} R_k(\boldsymbol{\theta}, \mathbf{r}).$$

It will also be convenient to work with approximate versions of the term $R_k(\boldsymbol{\theta}, \mathbf{r})$ in (D.9). Hence, we define below we define $\mathbb{R}_k(\boldsymbol{\theta}, \mathbf{r})$ which is essentially obtained by Taylor-expanding the term $R_{-k}(\boldsymbol{\theta})$ in (D.9).

$$(D.13) \quad S_k(\boldsymbol{\theta}, \mathbf{r}) = \Phi_{-k} + \frac{1}{2} (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{H}_{-k} (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*) + \frac{1}{n} \ell(\boldsymbol{\theta}; \mathbf{r}, y_k),$$

where $\mathbf{H}_{-k}$ is the Hessian of $R_{-k}(\boldsymbol{\theta})$ at $\boldsymbol{\theta}_{-k}^*$, i.e.

$$(D.14) \quad \mathbf{H}_{-k} = \nabla^2 R_{-k}(\boldsymbol{\theta})|_{\boldsymbol{\theta}=\boldsymbol{\theta}_{-k}^*}.$$

Finally, we denote the minimizer of $S(\boldsymbol{\theta}, \mathbf{r})$ by

$$(D.15) \quad \tilde{\boldsymbol{\theta}}_k(\mathbf{r}) = \arg \min_{\boldsymbol{\theta}} S_k(\boldsymbol{\theta}, \mathbf{r}),$$

and

$$(D.16) \quad \Psi_k(\mathbf{r}) = \min_{\boldsymbol{\theta}} S_k(\boldsymbol{\theta}, \mathbf{r}).$$

Simplification of Notation. We note from (D.7) that in our analysis the feature vectors are either $\mathbf{r}_i = \mathbf{a}_i$ or $\mathbf{r}_i = \mathbf{b}_i$; i.e. we can write Φ_k as

$$(D.17) \quad \Phi_k = \min_{\boldsymbol{\theta}} \frac{1}{n} \sum_{i=1}^n \ell(\boldsymbol{\theta}; \mathbf{r}_i, y_i) + \lambda \|\boldsymbol{\theta}\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \lambda_s \frac{\log(d)}{d} (\mathbf{1}^\top \boldsymbol{\theta})^2,$$

where the feature vectors $\mathbf{r}_i$ they are generated according to one of the following distributions

$$(D.18) \quad \mathbf{r}_i = \sigma(\mathbf{W} \mathbf{x}_i) \quad \text{or} \quad \mathbf{r}_i = \mathbf{f}_i := \mu_1 \mathbf{W} \mathbf{x}_i + \mu_2 \mathbf{u}_i,$$

It will be sometimes easier in our analysis to use (D.17), i.e. use $\mathbf{r}_i$ for both $\sigma(\mathbf{W} \mathbf{x}_i)$ and $\mathbf{f}_i$, but we will keep in mind that for $i \leq k$ we have $\mathbf{r}_i = \mathbf{f}_i$ and for $i > k$ we have $\mathbf{r}_i = \sigma(\mathbf{W} \mathbf{x}_i)$.

Details of the Gradient and Hessian of ℓ . In the following, we will need to work out the first and second derivatives of the loss function ℓ , given in (D.6), at multiple points. In order to present the derivations more compactly, let us denote

$$(D.19) \quad h(\boldsymbol{\theta}; \mathbf{r}, y) := (\boldsymbol{\theta}^\top \mathbf{r} - y)^2 \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2,$$

and provide the details for the derivatives of the loss function ℓ here. Given how the function h is defined, we can write

$$\ell(\boldsymbol{\theta}; \mathbf{r}, y) = (\boldsymbol{\theta}^\top \mathbf{r} - y)^2 + \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2 + 2g(h(\boldsymbol{\theta}; \mathbf{r}, y))$$

Using the notation ∇ for gradient w.r.t. $\boldsymbol{\theta}$, and ∇^2 for hessian w.r.t. $\boldsymbol{\theta}$, we can write

$$(D.20) \quad \nabla \ell(\boldsymbol{\theta}; \mathbf{r}, y) = 2(\boldsymbol{\theta}^\top \mathbf{r} - y)\mathbf{r} + 2\mathbf{J}^\top \mathbf{J}\boldsymbol{\theta} + 2\nabla h(\boldsymbol{\theta}; \mathbf{r}, y)g'(h(\boldsymbol{\theta}; \mathbf{r}, y))$$

and

$$(D.21)$$

$$\nabla^2 \ell(\boldsymbol{\theta}; \mathbf{r}, y) = 2\left(\mathbf{r}\mathbf{r}^\top + \mathbf{J}^\top \mathbf{J} + \nabla^2 h(\boldsymbol{\theta}; \mathbf{r}, y)g'(h(\boldsymbol{\theta}; \mathbf{r}, y)) + \nabla h(\boldsymbol{\theta}; \mathbf{r}, y)(\nabla h(\boldsymbol{\theta}; \mathbf{r}, y))^\top g''(h(\boldsymbol{\theta}; \mathbf{r}, y))\right),$$

where

$$(D.22) \quad \nabla h(\boldsymbol{\theta}; \mathbf{r}, y) = 2\left((\boldsymbol{\theta}^\top \mathbf{r} - y)\|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2 \mathbf{r} + (\boldsymbol{\theta}^\top \mathbf{r} - y)^2 \mathbf{J}^\top \mathbf{J}\boldsymbol{\theta}\right),$$

and

$$(D.23)$$

$$\nabla^2 h(\boldsymbol{\theta}; \mathbf{r}, y) = 2\left(\|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2 \mathbf{r}\mathbf{r}^\top + 2(\boldsymbol{\theta}^\top \mathbf{r} - y)\mathbf{r}\boldsymbol{\theta}^\top \mathbf{J}^\top \mathbf{J} + 2(\boldsymbol{\theta}^\top \mathbf{r} - y)\mathbf{J}^\top \mathbf{J}\boldsymbol{\theta}\mathbf{r}^\top + (\boldsymbol{\theta}^\top \mathbf{r} - y)^2 \mathbf{J}^\top \mathbf{J}\right).$$

D.2.1. Some properties of the minimizers in (D.10) In this section, we will analyze some of the properties of the vector $\boldsymbol{\theta}_{-k}^*$ and its relation with $\tilde{\boldsymbol{\theta}}_{-k}(\mathbf{r})$, for $k \in [n]$. We first show some basic properties of the vectors $\boldsymbol{\theta}_{-k}^*$ and $\boldsymbol{\theta}_k^*(\mathbf{r})$.

Lemma D.1 Fix $k \in [n]$. The following hold with absolute constants $c, C > 0$:

(a) The vector $\boldsymbol{\theta}_{-k}^*$ is bounded in the ℓ_2 norm:

$$(D.24) \quad \mathbb{P}\left(\|\boldsymbol{\theta}_{-k}^*\|_{\ell_2} \geq v + C\right) \leq c \exp(-nv^2/c).$$

(b) The vector $\boldsymbol{\theta}_k^*(\mathbf{r})$ is bounded in the ℓ_2 norm:

$$(D.25) \quad \mathbb{P}\left(\|\boldsymbol{\theta}_k^*(\mathbf{r})\|_{\ell_2} \geq v + C\right) \leq c \exp(-nv^2/c).$$

(c) For an independently generated vector $\mathbf{r}$ we have

$$(D.26) \quad \mathbb{P}\left(|\mathbf{r}^\top \boldsymbol{\theta}_{-k}^*(\mathbf{r})| \geq v\right) \leq c \exp(-v/c).$$

(d) We also have

$$(D.27) \quad |\mathbf{1}^\top \boldsymbol{\theta}_{-k}^*| \leq C \sqrt{\frac{d}{\log(d)}},$$

with probability at least $1 - ce^{-cn}$.

The proof of this lemma is provided in Section D.2.7. We now show that the distance between the two minimizers $\boldsymbol{\theta}_{-k}^*$ and $\tilde{\boldsymbol{\theta}}_k(\mathbf{r})$ is of order $O(1/\sqrt{d})$.

Lemma D.2 Fix $k \in [n]$. Assuming that $\mathbf{r}$ is generated independently from $\boldsymbol{\theta}_{-k}^*$ and according to one of the distributions in (D.18). Then, there exist absolute constants $c, c' > 0$ such that

$$(D.28) \quad \mathbb{P}\left(\|\tilde{\boldsymbol{\theta}}_{-k}(\mathbf{r}) - \boldsymbol{\theta}_{-k}^*\|_{\ell_2} \geq \frac{v}{\sqrt{d}}\right) \leq c \exp(-v^{c'}/c).$$

Proof We start by noting that since $\tilde{\boldsymbol{\theta}}_k(\mathbf{r})$ is the minimizer of (D.13):

$$(D.29) \quad \tilde{\boldsymbol{\theta}}_k(\mathbf{r}) = \arg \min_{\boldsymbol{\theta}} \left\{ S(\boldsymbol{\theta}) := \frac{1}{2}(\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{H}_{-k}(\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*) + \frac{1}{n}\ell(\boldsymbol{\theta}; \mathbf{r}, y_k) \right\}$$

Observe that (i) the function S is λ -strongly convex due to the fact that $R_{-k}(\boldsymbol{\theta})$ is strongly-convex, and thus its Hessian $\mathbf{H}_{-k}$ is a PSD matrix with smallest eigenvalue lower-bounded by λ ; (ii) $S(\boldsymbol{\theta}) \geq 0$ for any $\boldsymbol{\theta}$, as $\mathbf{H}_{-k}$ is a PSD matrix and ℓ is always positive-valued. As a result, we can write

$$(D.30) \quad \|\tilde{\boldsymbol{\theta}}_k(\mathbf{r}) - \boldsymbol{\theta}_{-k}^*\|_{\ell_2}^2 \leq \frac{1}{\lambda} S(\boldsymbol{\theta}_{-k}^*)$$

We can then write from (D.6) that

$$\begin{aligned} S(\boldsymbol{\theta}_{-k}^*) &= \frac{1}{n} \ell(\boldsymbol{\theta}_{-k}^*; \mathbf{r}, y_k) \\ &\leq \frac{1}{n} C \max\{1, (\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y)^2, \|\mathbf{J}\boldsymbol{\theta}_{-k}^*\|_{\ell_2}^2\}, \end{aligned}$$

where $C > 0$ is an absolute constant. The proof now follows from the result of Lemma D.1 and the fact that $\|\mathbf{J}\|$ is bounded, as well as the fact that d and n grow in proportion to each other. $\blacksquare$

Given the above lemma, we can analyze the behavior of $\Psi_k(\mathbf{r})$, defined in (D.16), in more detail.

Lemma D.3 Fix $k \in [n]$. We have

$$(D.31) \quad \Psi_k(\mathbf{r}) = \Phi_{-k} + \frac{1}{n} \min_{\tau_1} \left\{ \frac{1}{2n} \left(\frac{\partial \tilde{\ell}(\tau_1, 0)}{\partial \tau_1} \mathbf{r} + \frac{\partial \tilde{\ell}(\tau_1, 0)}{\partial \tau_2} \mathbf{p} \right)^\top \mathbf{H}_{-k}^{-1} \left(\frac{\partial \tilde{\ell}(\tau_1, 0)}{\partial \tau_1} \mathbf{r} + \frac{\partial \tilde{\ell}(\tau_1, 0)}{\partial \tau_2} \mathbf{p} \right) + \tilde{\ell}(\tau_1, 0) \right\} + e,$$

where (i) we have $\mathbf{p}^\top = 2\boldsymbol{\theta}_{-k}^{*\top} \mathbf{J}^\top \mathbf{J}$; (ii) the function $\tilde{\ell}(\tau_1, \tau_2)$ is defined as

$$(D.32) \quad \tilde{\ell}(\tau_1, \tau_2) := \rho_1 + \rho_2 + \rho_3 \tau_1 + \tau_1^2 + \tau_2 + g((\rho_1 + \rho_3 \tau_1 + \tau_1^2)(\rho_2 + \tau_2)),$$

with $\rho_1 = (\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y_k)^2$, $\rho_2 = \|\mathbf{J}\boldsymbol{\theta}_{-k}^*\|_{\ell_2}^2$, and $\rho_3 = 2(\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y_k)$; and (iii) the value e satisfies

$$\mathbb{P}(|e| \geq \frac{v}{d^{\frac{3}{2}}}) \leq c \exp(-v^{c'}/c),$$

for absolute constants $c, c' > 0$.

Furthermore, assuming that τ_1^* is the minimizer of the optimization problem in (D.31), we have

$$(D.33) \quad \tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^* = \frac{1}{n} (\beta_1 \mathbf{H}_{-k}^{-1} \mathbf{r} + \beta_2 \mathbf{H}_{-k}^{-1} \mathbf{p}) + \mathbf{e},$$

where β_1, β_2 depend only on τ_1^* , as well as $\mathbf{r}^\top \boldsymbol{\theta}_{-k}^*$, and $\|\mathbf{J}\boldsymbol{\theta}_{-k}^*\|_{\ell_2}$; and

$$(D.34) \quad \mathbb{P}(\max\{d^{\frac{3}{2}} \|\mathbf{e}\|_{\ell_2}, |\beta_1|, |\beta_2|, \|\mathbf{p}\|_{\ell_2}\} \geq v) \leq c \exp(-v^{c'}/c).$$

Proof In the following, to simplify notation, we use $\tilde{\boldsymbol{\theta}}$ instead of $\tilde{\boldsymbol{\theta}}_k(\mathbf{r})$. We have from (D.13) and (D.16) that

$$(D.35) \quad \Psi_k(\mathbf{r}) = \Phi_{-k} + \min_{\boldsymbol{\theta}} \left\{ \frac{1}{2}(\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{H}_{-k}(\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*) + \frac{1}{n}\ell(\boldsymbol{\theta}; \mathbf{r}, y_k) \right\}.$$

We now decompose the term $\frac{1}{n}\ell(\boldsymbol{\theta}; \mathbf{r}, y_k)$ (see (D.6)) according to the following set of simple relations:

$$\begin{aligned}\boldsymbol{\theta}^\top \mathbf{r} - y &= \boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y + (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{r}, \\ \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2 &= \|\mathbf{J}\boldsymbol{\theta}_{-k}^*\|_{\ell_2}^2 + 2\boldsymbol{\theta}_{-k}^{*\top} \mathbf{J}^\top \mathbf{J}(\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*) + (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{J}^\top \mathbf{J}(\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*), \\ (\boldsymbol{\theta}^\top \mathbf{r} - y)^2 &= (\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y)^2 + 2(\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y)\mathbf{r}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*) + (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{r} \mathbf{r}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*),\end{aligned}$$

As a result, it is easy to obtain the following:

$$\begin{aligned}\frac{1}{n}\ell(\boldsymbol{\theta}; \mathbf{r}, y_k) &= \frac{1}{n} \left\{ \rho_1 + \rho_2 + \rho_3 (\mathbf{r}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)) + (\mathbf{r}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*))^2 + \mathbf{p}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*) \right\} \\ &\quad + \frac{1}{n} \|\mathbf{J}(\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)\|_{\ell_2}^2 \\ &\quad + \frac{1}{n} g \left((\rho_1 + \rho_3 (\mathbf{r}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)) + (\mathbf{r}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*))^2) (\rho_2 + \mathbf{p}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*) + \|\mathbf{J}(\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)\|_{\ell_2}^2) \right),\end{aligned}$$

where the parameters ρ_j , $j = 1, 2, 3$, and the vector $\mathbf{p}$ are defined in the following:

$$\begin{aligned}\text{(D.36)} \quad \rho_1 &= (\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y_k)^2, \\ \rho_2 &= \|\mathbf{J}\boldsymbol{\theta}_{-k}^*\|_{\ell_2}^2, \\ \rho_3 &= 2(\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y_k), \\ \mathbf{p}^\top &= 2\boldsymbol{\theta}_{-k}^{*\top} \mathbf{J}^\top \mathbf{J}.\end{aligned}$$

We note that none of these quantities depend on the optimization variable $\boldsymbol{\theta}$ and hence can be considered as constants w.r.t. the minimization procedure in (D.35).

It will be convenient to consider the following variables:

$$\text{(D.37)} \quad \tau_1 = \mathbf{r}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*), \quad \text{and} \quad \tau_2 = \mathbf{p}^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*), \quad \text{and} \quad \tau_3 = \|\mathbf{J}(\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)\|_{\ell_2}^2.$$

As a result, we can write

$$\begin{aligned}\frac{1}{n}\ell(\boldsymbol{\theta}; \mathbf{r}, y_k) &= \frac{1}{n} \left\{ \rho_1 + \rho_2 + \rho_3 \tau_1 + \tau_1^2 + \tau_2 + \tau_3 + g \left((\rho_1 + \rho_3 \tau_1 + \tau_1^2) (\rho_2 + \tau_2 + \tau_3) \right) \right\} \\ \text{(D.38)} \quad &:= \frac{1}{n} \tilde{\ell}(\tau_1, \tau_2, \tau_3).\end{aligned}$$

Now, from (D.35), we can write the following equation for $\tilde{\boldsymbol{\theta}}$ (as it is the minimizer):

$$\text{(D.39)} \quad \mathbf{H}_{-k}(\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) = -\frac{1}{n} \left(\frac{\partial \tilde{\ell}}{\partial \tau_1} \mathbf{r} + \frac{\partial \tilde{\ell}}{\partial \tau_2} \mathbf{p} + \frac{\partial \tilde{\ell}}{\partial \tau_3} \mathbf{J}^\top \mathbf{J}(\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \right),$$

where the partial derivatives are evaluated at τ_1, τ_2, τ_3 when $\boldsymbol{\theta} = \tilde{\boldsymbol{\theta}}$. Consequently, we have

$$\text{(D.40)} \quad \tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^* = -\frac{1}{n} \left(\frac{\partial \tilde{\ell}}{\partial \tau_1} \mathbf{H}_{-k}^{-1} \mathbf{r} + \frac{\partial \tilde{\ell}}{\partial \tau_2} \mathbf{H}_{-k}^{-1} \mathbf{p} + \frac{\partial \tilde{\ell}}{\partial \tau_3} \mathbf{H}_{-k}^{-1} \mathbf{J}^\top \mathbf{J}(\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \right),$$

Using the above relation, we can derive a few useful properties. First, given how ρ_j 's are defined in (D.36), and by using Lemma D.1, and since $g(x) = \sqrt{x + \gamma}$ has uniformly bounded first and second derivative, and by using the fact that the operator norm of $\mathbf{J}$ is bounded, it is easy to show that for absolute constants c, c' we have

$$\text{(D.41)} \quad \mathbb{P} \left(\max \left\{ \left| \frac{\partial \tilde{\ell}}{\partial \tau_1} \right|, \left| \frac{\partial \tilde{\ell}}{\partial \tau_2} \right|, \left| \frac{\partial \tilde{\ell}}{\partial \tau_3} \right| \right\} \geq v \right) \leq c \exp(-v^{c'}/c),$$

where in the above relation the partial derivatives are evaluated at τ_1, τ_2, τ_3 , when $\boldsymbol{\theta} = \tilde{\boldsymbol{\theta}}$.

Second, by using Lemma D.2, and the fact that the norm of the matrix $\mathbf{J}$ is bounded, as well as the fact that the operator norm of $\mathbf{H}_{-k}^{-1}$ is upper-bounded by $1/\lambda$, we have for absolute constants c, c' that

$$(D.42) \quad \mathbb{P} \left(\left\| \frac{1}{n} \mathbf{H}_{-k}^{-1} \mathbf{J}^\top \mathbf{J} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \right\|_{\ell_2} \geq \frac{v}{d^{\frac{3}{2}}} \right) \leq c \exp(-v^{c'}/c).$$

Third, from the definition of $\mathbf{p}$ in (D.36), and by using Lemma D.1 as well as (D.40) we obtain for absolute constants c, c' that

$$(D.43) \quad \mathbb{P} \left(\mathbf{p}^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \geq \frac{v}{d^{\frac{1}{2}}} \right) \leq c \exp(-v^{c'}/c).$$

The above relation shows that the value of τ_2 , evaluated at $\boldsymbol{\theta} = \tilde{\boldsymbol{\theta}}$, is of order $O(d^{-\frac{1}{2}})$.

Fourth, we can write using Lemma D.2 that

$$(D.44) \quad \mathbb{P} \left(\left\| \mathbf{J} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \right\|_{\ell_2} \geq \frac{v}{d^{\frac{1}{2}}} \right) \leq c \exp(-v^{c'}/c),$$

which essentially results in τ_3 , evaluated at $\boldsymbol{\theta} = \tilde{\boldsymbol{\theta}}$, to be of the order $O(d^{-1})$.

Finally, we note that for each of the partial derivatives we can write

$$(D.45) \quad \mathbb{P} \left(\left| \frac{\partial \tilde{\ell}(\tau_1, \tau_2, \tau_3)}{\partial \tau_j} - \frac{\partial \tilde{\ell}(\tau_1, 0, 0)}{\partial \tau_j} \right| \geq \frac{v}{d^{\frac{1}{2}}} \right) \leq c \exp(-v^{c'}/c),$$

for $j = 1, 2, 3$ and for τ_j 's that are evaluated at $\boldsymbol{\theta} = \tilde{\boldsymbol{\theta}}$.

Using the above five properties, we can conclude that

$$(D.46) \quad \tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^* = -\frac{1}{n} \left(\frac{\partial \tilde{\ell}(\tau_1, 0, 0)}{\partial \tau_1} \mathbf{H}_{-k}^{-1} \mathbf{r} + \frac{\partial \tilde{\ell}(\tau_1, 0, 0)}{\partial \tau_2} \mathbf{H}_{-k}^{-1} \mathbf{p} \right) + \mathbf{e},$$

and

$$(D.47) \quad \frac{1}{n} \tilde{\ell}(\tau_1, \tau_2, \tau_3) = \frac{1}{n} \tilde{\ell}(\tau_1, 0, 0) + e',$$

where τ_1, τ_2, τ_3 are computed from (D.37) at $\boldsymbol{\theta} = \tilde{\boldsymbol{\theta}}$, and

$$\mathbb{P} \left(\max\{\|\mathbf{e}\|_{\ell_2}, |e'|\} \geq \frac{v}{d^{\frac{3}{2}}} \right) \leq c \exp(-v^{c'}/c).$$

As a result, by defining

$$\tilde{\ell}(\tau_1, \tau_2) := \tilde{\ell}(\tau_1, \tau_2, 0),$$

where $\tilde{\ell}$ is given in (D.38), and by plugging the solution (D.46) into the optimization in (D.35), we obtain

$$(D.48) \quad \Psi_k(\mathbf{r}) = \Phi_{-k} + \frac{1}{n} \min_{\tau_1} \left\{ \frac{1}{2n} \left(\frac{\partial \tilde{\ell}(\tau_1, 0)}{\partial \tau_1} \mathbf{r} + \frac{\partial \tilde{\ell}(\tau_1, 0)}{\partial \tau_2} \mathbf{p} \right)^\top \mathbf{H}_{-k}^{-1} \left(\frac{\partial \tilde{\ell}(\tau_1, 0)}{\partial \tau_1} \mathbf{r} + \frac{\partial \tilde{\ell}(\tau_1, 0)}{\partial \tau_2} \mathbf{p} \right) + \tilde{\ell}(\tau_1, 0) \right\} + e,$$

where

$$\mathbb{P} \left(|e| \geq \frac{v}{d^{\frac{3}{2}}} \right) \leq c \exp(-v^{c'}/c).$$

■

D.2.2. Bounding $\|\tilde{\boldsymbol{\theta}}(\mathbf{r}) - \boldsymbol{\theta}^*(\mathbf{r})\|_{\ell_2}$

Lemma D.4 Fix $k \in [n]$. There exist absolute constants $b, c, c' > 0$ such that for any $v \geq 0$

$$(D.49) \quad \mathbb{P}\left(\|\tilde{\boldsymbol{\theta}}_k(\mathbf{r}) - \boldsymbol{\theta}_k^*(\mathbf{r})\|_{\ell_2} \geq v \frac{(\log(d))^b}{d}\right) \leq c \exp\left(-\frac{vc'}{c}\right) + c \exp\left(-(\log(d))^2/c\right)$$

Proof To simplify notation, in this lemma we use $\boldsymbol{\theta}^*$ instead of $\boldsymbol{\theta}_k^*(\mathbf{r})$ and $\tilde{\boldsymbol{\theta}}$ instead of $\tilde{\boldsymbol{\theta}}_k(\mathbf{r})$. Also, we define

$$q(\boldsymbol{\theta}) = \lambda \|\boldsymbol{\theta}\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \lambda_s \frac{\log(d)}{d} (\mathbf{1}^\top \boldsymbol{\theta})^2$$

Due to the fact that $R(\boldsymbol{\theta}, \mathbf{r})$ is λ -strongly convex, we can write

$$(D.50) \quad \|\boldsymbol{\theta}^* - \tilde{\boldsymbol{\theta}}\|_{\ell_2} \leq \frac{1}{\lambda} \|\nabla_{\boldsymbol{\theta}} R(\tilde{\boldsymbol{\theta}}, \mathbf{r})\|_{\ell_2}$$

Consequently, we will bound the right-hand-side in the above relation. By using the fact that $\boldsymbol{\theta}_{-k}^*$ is the minimizer of the function R_{-k} , we can write

$$(D.51) \quad \nabla_{\boldsymbol{\theta}} R(\tilde{\boldsymbol{\theta}}, \mathbf{r}) = \nabla_{\boldsymbol{\theta}} R(\tilde{\boldsymbol{\theta}}, \mathbf{r}) - \nabla_{\boldsymbol{\theta}} R_{-k}(\boldsymbol{\theta}_{-k}^*).$$

From (D.9) and (D.51) we have

$$\begin{aligned} \nabla_{\boldsymbol{\theta}} R(\tilde{\boldsymbol{\theta}}, \mathbf{r}) &= \nabla R_{-k}(\tilde{\boldsymbol{\theta}}) + \frac{1}{n} \nabla \ell(\tilde{\boldsymbol{\theta}}; \mathbf{r}, y_k) - \nabla R_{-k}(\boldsymbol{\theta}_{-k}^*) \\ &= \frac{1}{n} \sum_{t \neq k} (\nabla \ell(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) - \nabla \ell(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) + \nabla q(\tilde{\boldsymbol{\theta}}) - \nabla q(\boldsymbol{\theta}_{-k}^*) + \frac{1}{n} \nabla \ell(\tilde{\boldsymbol{\theta}}; \mathbf{r}, y_k), \end{aligned}$$

By using the fact that $\tilde{\boldsymbol{\theta}}$ is the minimizer of (D.13), we can write

$$\nabla_{\boldsymbol{\theta}} R(\tilde{\boldsymbol{\theta}}, \mathbf{r}) = \frac{1}{n} \sum_{t \neq k} (\nabla \ell(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) - \nabla \ell(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) - \mathbf{H}_{-k}(\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) + \nabla q(\tilde{\boldsymbol{\theta}}) - \nabla q(\boldsymbol{\theta}_{-k}^*)$$

Now, from (D.14), we obtain

$$(D.52) \quad \nabla_{\boldsymbol{\theta}} R(\tilde{\boldsymbol{\theta}}, \mathbf{r}) = \frac{1}{n} \sum_{t \neq k} \nabla \ell(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) - \nabla \ell(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t) - \nabla^2 \ell(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t) (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*),$$

where in the above we have used the fact that, since q is a quadratic function, we have $\nabla q(\tilde{\boldsymbol{\theta}}) - \nabla q(\boldsymbol{\theta}_{-k}^*) - \nabla^2 q(\boldsymbol{\theta}_{-k}^*) (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) = 0$. We will now analyze each of the terms above. We first bound the first term (i.e. the sum involving the derivatives of ℓ). We will use the following simple relations for any choice of $\mathbf{r}, y$:

$$\begin{aligned} \tilde{\boldsymbol{\theta}}^\top \mathbf{r} - y &= \boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y + (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{r}, \\ \|J\tilde{\boldsymbol{\theta}}\|_{\ell_2}^2 &= \|J\boldsymbol{\theta}_{-k}^*\|_{\ell_2}^2 + 2\boldsymbol{\theta}_{-k}^{*\top} J^\top J (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) + (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top J^\top J (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*), \\ (\tilde{\boldsymbol{\theta}}^\top \mathbf{r} - y)^2 &= (\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y)^2 + 2(\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y) \mathbf{r}^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) + (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{r} \mathbf{r}^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*), \end{aligned}$$

and

$$\begin{aligned} h(\tilde{\boldsymbol{\theta}}; \mathbf{r}, y) - h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}, y) &= (\tilde{\boldsymbol{\theta}}^\top \mathbf{r} - y)^2 \|J\tilde{\boldsymbol{\theta}}\|_{\ell_2}^2 - (\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y)^2 \|J\boldsymbol{\theta}_{-k}^*\|_{\ell_2}^2 \\ &= \nabla h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}, y)^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) + e(\mathbf{r}, y) \\ &= 2(\|J\boldsymbol{\theta}_{-k}^*\|_{\ell_2}^2 \mathbf{r}^\top + (\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y)^2 J^\top J \boldsymbol{\theta}_{-k}^{*\top}) (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) + e(\mathbf{r}, y), \end{aligned}$$

where the error term $e(\mathbf{r}, y)$ can be written as

$$e(\mathbf{r}, y) = (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \nabla_{\boldsymbol{\theta}}^2 h(\boldsymbol{\theta}; \mathbf{r}, y) |_{\boldsymbol{\theta}=\boldsymbol{\theta}(\mathbf{r}, y)} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*),$$

and $\boldsymbol{\theta}(\mathbf{r}, y) = \zeta \boldsymbol{\theta}_{-k}^* + (1 - \zeta) \tilde{\boldsymbol{\theta}}$ for some $\zeta \in [0, 1]$ which depends on $\mathbf{r}$ and y .

We will also use the Taylor expansion:

$$\begin{aligned} g'(h(\tilde{\boldsymbol{\theta}}; \mathbf{r}, y_t)) &= g'(h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) + g''(h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) (h(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) - h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) \\ &\quad + \frac{1}{2} g'''(v_t) (h(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) - h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t))^2, \end{aligned}$$

where v_t is a number between $h(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t)$ and $h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)$.

From (D.20) and (D.21) we will decompose:

$$\nabla \ell(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) - \nabla \ell(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t) - \nabla^2 \ell(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t) (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) = \text{Term}_1 + \text{Term}_2 - \text{Term}_3,$$

where the terms are given in (D.53), (D.55), and (D.56). We will bound each of these terms in the following. We have

$$(D.53) \quad \text{Term}_1 = 2 \left(\nabla h(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) - \nabla h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t) \right) g'(h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t))$$

where

$$\begin{aligned} &\nabla h(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) - \nabla h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t) \\ &= 2 \left((\tilde{\boldsymbol{\theta}}^\top \mathbf{r}_t - y_t) \left\| \mathbf{J} \tilde{\boldsymbol{\theta}} \right\|_{\ell_2}^2 \mathbf{r}_t + (\tilde{\boldsymbol{\theta}}^\top \mathbf{r}_t - y_t)^2 \mathbf{J}^\top \mathbf{J} \tilde{\boldsymbol{\theta}} - (\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r}_t - y_t) \left\| \mathbf{J} \boldsymbol{\theta}_{-k}^* \right\|_{\ell_2}^2 \mathbf{r}_t + (\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r}_t - y_t)^2 \mathbf{J}^\top \mathbf{J} \boldsymbol{\theta}_{-k}^* \right) \end{aligned}$$

And thus,

$$\begin{aligned} (D.54) \quad &\nabla h(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) - \nabla h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t) \\ &= 2 \left((\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{r}_t \left\| \mathbf{J} \boldsymbol{\theta}_{-k}^* \right\|_{\ell_2}^2 + (\tilde{\boldsymbol{\theta}}^\top \mathbf{r}_t - y_t) \left(2 \boldsymbol{\theta}_{-k}^{*\top} \mathbf{J}^\top \mathbf{J} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) + (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{J}^\top \mathbf{J} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \right) \right) \mathbf{r}_t \\ &\quad + 2 \left((\tilde{\boldsymbol{\theta}}^\top \mathbf{r}_t - y_t)^2 \mathbf{J}^\top \mathbf{J} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) + (2 (\tilde{\boldsymbol{\theta}}^\top \mathbf{r}_t - y_t) \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) + (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{r}_t \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)) \mathbf{J}^\top \mathbf{J} \tilde{\boldsymbol{\theta}} \right) \end{aligned}$$

We also have

$$\begin{aligned} (D.55) \quad \text{Term}_2 &= 2 \nabla h(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) g''(h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) (h(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) - h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) \\ &= 2 \nabla h(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) g''(h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) (\nabla h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) + e(\mathbf{r}_t, y_t)), \end{aligned}$$

and

$$\begin{aligned} (D.56) \quad \text{Term}_3 &= 2 \left(\nabla^2 h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t) g'(h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) + \nabla h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t) (\nabla h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t))^\top g''(h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) \right) (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \\ &\quad - \frac{1}{2} g'''(v_t) (h(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) - h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t))^2, \end{aligned}$$

After some straight-forward steps, we can write

$$(D.57)$$

$$\text{Term}_1 + \text{Term}_2 - \text{Term}_3$$

$$\begin{aligned} &= 4 g'(h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) \left(2 \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \boldsymbol{\theta}_{-k}^{*\top} \mathbf{J}^\top \mathbf{J} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \mathbf{r}_t \right. \\ &\quad \left. + 2 \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \left((\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r}_t - y_t) \mathbf{J}^\top \mathbf{J} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) + \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \mathbf{J}^\top \mathbf{J} \tilde{\boldsymbol{\theta}} \right) \right) \end{aligned}$$

$$\begin{aligned}
& + \left(2(\boldsymbol{\theta}_{-k}^*{}^\top \mathbf{r}_t - y_t) \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) + (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{r}_t \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \right) \mathbf{J}^\top \mathbf{J} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \\
& + (\mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*))^2 \mathbf{J}^\top \mathbf{J} \tilde{\boldsymbol{\theta}} + \|\mathbf{J}(\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)\|_{\ell_2}^2 \mathbf{r}_t \\
& + 2g''(h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) \left((\nabla h(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) - \nabla h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \nabla h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t) + \nabla h(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) e(\mathbf{r}_t, y_t) \right) \\
& - \frac{1}{2} g'''(v_t) (h(\tilde{\boldsymbol{\theta}}; \mathbf{r}_t, y_t) - h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t))^2.
\end{aligned}$$

The relation (D.57) has itself three different terms. We will now simplify and bound each of the terms above. However, we remark that all the three terms will be bounded in a similar way. Let's consider the first term in the right-hand-side of (D.57). The first part of this term is:

$$4g'(h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) \times 2\mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \boldsymbol{\theta}_{-k}^*{}^\top \mathbf{J}^\top \mathbf{J} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \mathbf{r}_t,$$

which can be rewritten as

$$8g'(h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) \boldsymbol{\theta}_{-k}^*{}^\top \mathbf{J}^\top \mathbf{J} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \mathbf{r}_t \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*).$$

Now, by using the fact that the first derivatives of the function g is uniformly bounded, and by some straight-forward usages of the Cauchy-Schwartz inequality, we can easily rewrite the above part as

$$\alpha_{1,t} \mathbf{p}^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \mathbf{r}_t \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*),$$

where the vector $\mathbf{p} = \boldsymbol{\theta}_{-k}^*{}^\top \mathbf{J}^\top \mathbf{J}$ does not depend on t , and $\alpha_{1,t}$ is a constant. We can further write:

$$|\alpha_{1,t}| \leq C \text{ and } \|\mathbf{p}\|_{\ell_2} \leq \|\mathbf{J}\|_{\ell_2}^2 \|\boldsymbol{\theta}_{-k}^*\|_{\ell_2} \leq \|\mathbf{J}\|_{\ell_2}^4 + \|\boldsymbol{\theta}_{-k}^*\|_{\ell_2}^2,$$

where C is an absolute constant.

Let us now consider the second part of the first term in (D.57), which is:

$$4g'(h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) \times 2\mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) (\boldsymbol{\theta}_{-k}^*{}^\top \mathbf{r}_t - y_t) \mathbf{J}^\top \mathbf{J} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*).$$

We can rewrite this part as

$$8g'(h(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) \times (\boldsymbol{\theta}_{-k}^*{}^\top \mathbf{r}_t - y_t) \mathbf{J}^\top \mathbf{J} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{r}_t (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) = \alpha_{2,t} \mathbf{A} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{r}_t (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*),$$

where, the matrix $\mathbf{A}$ is the same for all t , and

$$|\alpha_{2,t}| \leq C |\boldsymbol{\theta}_{-k}^*{}^\top \mathbf{r}_t - y_t| \text{ and } \|\mathbf{A}\|_{\ell_2} \leq \|\mathbf{J}\|_{\ell_2}^2.$$

In a similar way, one can inspect all the parts of the first term in (D.57) and show that they take the form of one of the following:

$$\begin{aligned}
& \bullet \alpha_{1,t} \mathbf{p}_1^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \mathbf{r}_t \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*), \\
& \bullet \alpha_{2,t} \mathbf{A}_1 (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{r}_t (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*), \\
& \bullet \alpha_{3,t} \mathbf{p}_2 (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{r}_t \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*), \\
& \bullet \alpha_{4,t} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{A}_2 (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \mathbf{r}_t, \\
& \bullet \alpha_{5,t} ((\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{r}_t)^2 \mathbf{p}_3,
\end{aligned}
\tag{D.58}$$

where $\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_3, \mathbf{A}_1, \mathbf{A}_2$ do not depend on t , and we have

(D.59)

$$|\alpha_{j,t}| \leq C(1 + (\boldsymbol{\theta}_{-k}^*{}^\top \mathbf{r}_t - y_t)^2) \text{ and } \max\{|\gamma|, \|\mathbf{p}_1\|_{\ell_2}, \|\mathbf{p}_2\|_{\ell_2}, \|\mathbf{p}_3\|_{\ell_2}, \|\mathbf{A}_1\|_{\ell_2}, \|\mathbf{A}_2\|_{\ell_2}\} \leq C(1 + \|\mathbf{J}\|_{\ell_2}^2).$$

We will now consider the second term in the right-hand-side of (D.57) which can be expanded using (D.22) and (D.54). Again, one can inspect all the parts and show that they take one of the forms in the following (in addition to the forms presented in (D.58)):

$$(D.60) \quad \begin{aligned} & \bullet \alpha_{3,t} (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \mathbf{r}_t \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{p}_3, \\ & \bullet \alpha_{4,t} (\mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*))^2 \mathbf{r}_t, \\ & \bullet \alpha_{5,t} \mathbf{p}_4^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \mathbf{r}_t \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*), \end{aligned}$$

where, for some positive constant C and even integer D we have

$$(D.61) \quad |\alpha_{j,t}| \leq C(1 + (\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y)^D + (\tilde{\boldsymbol{\theta}}^\top \mathbf{r} - y)^D) \text{ and } \max\{\|\mathbf{p}_3\|_{\ell_2}, \|\mathbf{p}_4\|_{\ell_2}\} \leq C(1 + \|\mathbf{J}\|_{\ell_2}^D + \|\boldsymbol{\theta}_{-k}^*\|_{\ell_2}^D).$$

A similar bounding can be done for the third term in (D.57).

We now claim that the sum of each of the terms in (D.58) and (D.60) over t is at most of order $O(\text{polylog}(n)/n)$. Let's consider the first term in (D.58). We can write

$$\begin{aligned} \frac{1}{n} \left\| \sum_t \alpha_{1,t} \mathbf{p}_1^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \mathbf{r}_t \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \right\|_{\ell_2} & \leq \frac{1}{n} \sup_t \{|\alpha_{1,t}|\} |\mathbf{p}_1^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)| \left\| \sum_{t \neq k} \mathbf{r}_t \mathbf{r}_t^\top \right\|_{\ell_2} \|\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*\|_{\ell_2} \\ & = \sup_t \{|\alpha_{1,t}|\} |\mathbf{p}_1^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)| \left\| \frac{1}{n} \sum_{t \neq k} \mathbf{r}_t \mathbf{r}_t^\top \right\|_{\ell_2} \|\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*\|_{\ell_2} \end{aligned}$$

Now, from Lemma D.5 it should be clear why the above quantity is small. Informally, and neglecting the polylog factors, the lemma asserts that with high probability the terms $\sup_t \{|\alpha_{1,t}|\}$, and $\left\| \frac{1}{n} \sum_{t \neq k} \mathbf{r}_t \mathbf{r}_t^\top \right\|_{\ell_2}$ are all $O(1)$; but the term $|\mathbf{p}_1^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)|$, is $O(1/n)$ as $\mathbf{p}_1$ is a fixed vector (independent of t), and $\|\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*\|_{\ell_2}$ is $O(n^{-\frac{1}{2}})$. As a result, the whole expression is $O(n^{-\frac{3}{2}})$. Formally, it is easy to conclude from Lemma D.5 that for a given $k \in [d]$:

$$(D.62) \quad \mathbb{P} \left(\frac{1}{n} \left\| \sum_t \alpha_{1,t} \mathbf{p}_1^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \mathbf{r}_t \mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \right\|_{\ell_2} \geq v \frac{(\log(d))^b}{d^{\frac{3}{2}}} \right) \leq c \exp(-v^{c'}/c) + c \exp(-(\log(d))^2/c),$$

for some absolute constants $c, c', c'' > 0$.

Let's now consider the second term in (D.58). We can write

$$\begin{aligned} \left\| \frac{1}{n} \sum_t \alpha_{2,t} \mathbf{A}_1 (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \mathbf{r}_t (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*) \right\|_{\ell_2} & = \|\mathbf{A}_1\|_{\ell_2} \|\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*\|_{\ell_2} \left\| (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \frac{1}{n} \sum_{t \neq k} \alpha_{2,t} \mathbf{r}_t \right\|_{\ell_2} \\ & = \|\mathbf{A}_1\|_{\ell_2} \|\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*\|_{\ell_2} \left\| (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \times \frac{1}{n} \sum_{t \neq k} \alpha_{2,t} \mathbf{r}_t \right\|_{\ell_2}. \end{aligned}$$

Now, note from the first part of Lemma D.5 that the norm of the vector $\frac{1}{n} \sum_{t \neq k} \alpha_{2,t} \mathbf{r}_t^\top$ is w.h.p. $O(1)$. Also, this vector is independent from $\mathbf{r}$, and hence from the second part of Lemma D.5 we obtain that w.h.p. $\|(\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)^\top \times \frac{1}{n} \sum_{t \neq k} \alpha_{2,t} \mathbf{r}_t\|_{\ell_2}$ is $O(n^{-\frac{1}{2}})$. Consequently, by noting that $\|\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*\|_{\ell_2}$ is w.h.p. $O(n^{-\frac{1}{2}})$ we obtain that the whole expression is w.h.p. $O(n^{-\frac{3}{2}})$. The formal expression would be like the probabilistic expression given in (D.62).

Similarly, for the term $\alpha_{4,t} (\mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*))^2 \mathbf{r}_t$ we can write

$$(D.63) \quad \left\| \frac{1}{n} \sum_{t \neq k} \alpha_{4,t} (\mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*))^2 \mathbf{r}_t \right\|_{\ell_2} \leq \sup_t \{(\mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*))^2\} \left\| \frac{1}{n} \sum_{t \neq k} \alpha_{4,t} \mathbf{r}_t \right\|_{\ell_2}.$$

Now, by part (e) of Lemma D.5 we can easily conclude that

$$\mathbb{P}\left(\sup_t \{(\mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*))^2\} \geq (\log(d))^b\right) \leq c \exp(-(\log(d))^2/c),$$

for absolute constants $b, c > 0$ that are suitably chosen. A similar conclusion can be made for $\sup_t \{|\alpha_{4,t}|\}$ from (D.61) and part (g) of Lemma D.5–i.e.

$$\mathbb{P}\left(\sup_t \{|\alpha_{4,t}|\} \geq (\log(d))^b\right) \leq c \exp(-(\log(d))^2/c).$$

As a result, by using the above bounds, as well as (D.63), and part (a) of Lemma D.5 we obtain

$$\mathbb{P}\left(\left\|\frac{1}{n} \sum_{t \neq k} \alpha_{4,t} (\mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*))^2 \mathbf{r}_t\right\|_{\ell_2} \geq v \frac{(\log(d))^b}{d}\right) \leq c \exp(-v^{c'}/c) + c \exp(-(\log(d))^2/c),$$

In a similar way as the above, we can show that the sum of all the terms in (D.58) and (D.60) over t have similar bounds. As a result, going back to (D.52), we have shown that the sum can be bounded to give the desired result as in the lemma. ■

Lemma D.5 *For some absolute constants $b, c, c' > 0$ we have:*

(a)

$$\mathbb{P}\left(\left\|\frac{1}{n} \sum_{t \neq k} \mathbf{r}_t \mathbf{r}_t^\top\right\|_{\ell_2} \geq v\right) \leq c \exp(-v^{c'}/c).$$

(b) *Given any sequence of numbers $\{\alpha_t\}_{t=1}^n$ such that $|\alpha_t| \leq 1$, we have*

$$\mathbb{P}\left(\left\|\frac{1}{n} \sum_{t \neq k} \alpha_t \mathbf{r}_t\right\|_{\ell_2} \geq v\right) \leq c e^{-v^{c'}/c}.$$

(c) *Given a vector $\mathbf{u}$, with $\|\mathbf{u}\|_{\ell_2} = 1$, we can write*

$$\mathbb{P}(|\mathbf{u}^\top \mathbf{r}| \geq v) \leq c e^{-v/c}.$$

(d) *Given a vector $\mathbf{u}$, with $\|\mathbf{u}\|_{\ell_2} = 1$, which is independent from $\mathbf{r}$, we have*

$$\mathbb{P}\left(|\mathbf{u}^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)| \geq \frac{v}{n}\right) \leq c e^{-v/c}.$$

(e) *For $\mathbf{r}$ generated according to either of the distributions in (D.18), we have*

$$\mathbb{P}(\|\mathbf{r}\|_{\ell_2} \geq v\sqrt{n}) \leq c \exp(-v^2/c).$$

(f) *We further have*

$$\mathbb{P}\left(|\mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}} - \boldsymbol{\theta}_{-k}^*)| \geq \frac{v}{\sqrt{n}}\right) \leq c e^{-v/c}.$$

(g) *For a given even integer $D > 0$ we have:*

$$\mathbb{P}\left(\max\left\{(\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y)^D, (\tilde{\boldsymbol{\theta}}^\top \mathbf{r} - y)^D, \|\boldsymbol{\theta}_{-k}^*\|_{\ell_2}^D\right\} \geq v\right) \leq c e^{-v^{c'}/c}.$$

As a corollary, we have

$$\mathbb{P}\left(\max\left\{(\boldsymbol{\theta}_{-k}^{*\top} \mathbf{r} - y)^D, (\tilde{\boldsymbol{\theta}}^\top \mathbf{r} - y)^D, \|\boldsymbol{\theta}_{-k}^*\|_{\ell_2}^D\right\} \geq (\log(d))^b\right) \leq c \exp\{-(\log(d))^2/c\}.$$

Proof of this lemma is provided in Section D.2.7.

D.2.3. Bounding $|\mathbf{r}^\top \tilde{\boldsymbol{\theta}}(\mathbf{r})|$

Lemma D.6 *There exist absolute constants $c, c' > 0$ such that for every $k \in [n]$ and $v \geq 0$ we have*

$$(D.64) \quad \mathbb{P}(|\mathbf{r}^\top \tilde{\boldsymbol{\theta}}(\mathbf{r})| \geq v) \leq c \exp(-v^{c'}/c),$$

and

$$(D.65) \quad \mathbb{P}(|\mathbf{r}^\top \boldsymbol{\theta}_k^*(\mathbf{r})| \geq v) \leq cd \exp(-v^{c'}/c) + c \exp(-(\log(d))^2/c).$$

Proof This lemma can also be proven similarly to [39] (see Section D.3 in [39]). There are, however, some small differences that we will mention the details here.

For the first part, the proof proceeds in two steps. In this first step, we use Lemma D.1 to obtain

$$(D.66) \quad \mathbb{P}(|\mathbf{r}^\top \boldsymbol{\theta}_{-k}^*| \geq v) \leq c \exp(-v/c).$$

We will now bound the term $|\mathbf{r}^\top (\tilde{\boldsymbol{\theta}}(\mathbf{r}) - \boldsymbol{\theta}_{-k}^*(\mathbf{r}))|$. We can write:

$$\begin{aligned} \mathbb{P}(|\mathbf{r}^\top (\tilde{\boldsymbol{\theta}}(\mathbf{r}) - \boldsymbol{\theta}_{-k}^*(\mathbf{r}))| \geq v) &= \mathbb{P}\left(\left|\frac{1}{\sqrt{d}} \mathbf{r}^\top \times \sqrt{d}(\tilde{\boldsymbol{\theta}}(\mathbf{r}) - \boldsymbol{\theta}_{-k}^*(\mathbf{r}))\right| \geq v\right) \\ &\leq \mathbb{P}\left(\|\tilde{\boldsymbol{\theta}}(\mathbf{r}) - \boldsymbol{\theta}_{-k}^*\|_{\ell_2} \geq \sqrt{\frac{v}{d}}\right) + \mathbb{P}\left(\|\mathbf{r}\|_{\ell_2} \geq \sqrt{vd}\right) \end{aligned}$$

Now, the first term above can be bounded using Lemma D.2, and the second can be bounded from Lemma D.5, and thus the proof of the lemma follows.

The proof of the second part is similar to the first part (and we use Lemma D.4). $\blacksquare$

D.2.4. Bounding $\|\boldsymbol{\theta}_{-k}^*\|_{\ell_\infty}$

Lemma D.7 *Let $\boldsymbol{\theta}_{-k}^*$ be the minimizer of R_{-k} defined in (D.8). There exist absolute constants $c_0, c_1, c_\infty > 0$ such that for any $k \in [n]$:*

$$(D.67) \quad \mathbb{P}\left\{\|\boldsymbol{\theta}_{-k}^*\|_{\ell_\infty} \geq c_\infty \sqrt{\frac{\log(d)}{d}}\right\} \leq 5d^{-c_0} + 3e^{-c_1 n}$$

Proof For convenience we remind the definition of $\boldsymbol{\theta}_{-k}^*$ given by $\boldsymbol{\theta}_{-k}^* = \arg \min_{\boldsymbol{\theta}} R_{-k}(\boldsymbol{\theta})$ where

$$(D.68) \quad R_{-k}(\boldsymbol{\theta}) = \frac{1}{n} \sum_{i=1}^{k-1} \ell(\boldsymbol{\theta}; \mathbf{a}_i, y_i) + \frac{1}{n} \sum_{i=k+1}^n \ell(\boldsymbol{\theta}; \mathbf{b}_i, y_i) + \lambda \|\boldsymbol{\theta}\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \lambda_s \frac{\sqrt{\log(d)}}{d} (\mathbf{1}^\top \boldsymbol{\theta})^2,$$

and

$$\ell(\boldsymbol{\theta}; \mathbf{r}, y) = (\boldsymbol{\theta}^\top \mathbf{r} - y)^2 + \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2 + 2g\left((\boldsymbol{\theta}^\top \mathbf{r} - y)^2 \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}^2\right),$$

with $g(x) = \sqrt{x + \gamma}$.

We use a similar strategy as in the proof of Proposition 6.2. Specifically we first bound the last coordinate of $\boldsymbol{\theta}_{-k}^*$. Next, by symmetry we conclude that the same bound holds for all of its coordinates and control its ℓ_∞ norm by union bounding.

With a slight abuse of notation, we consider a $(N+1)$ dimensional version of the above optimization over $[\boldsymbol{\theta}; u]$ and denote the last coordinate of the optimal solution by $\hat{u}$. We define

$$\mathbf{r}_i = \begin{cases} \sigma(\mathbf{W} \mathbf{x}_i) & \text{if } i \leq k-1, \\ \mathbf{0} & \text{if } i = k, \\ \mathbf{f}_i & \text{if } k \leq i \leq n. \end{cases}$$

Also define $\mathbf{e} = [a_1, \dots, a_{k-1}, b_{k+1}, \dots, b_n]$ with $a_i = \sigma(\mathbf{w}_{N+1}^\top \mathbf{x}_i)$ and $b_i = \mu_1 \mathbf{w}_{N+1}^\top \mathbf{x}_i + \mu_2 z_i$ (with $z_i \sim \mathcal{N}(0, 1)$ independent of $\mathbf{x}_i$). From (D.68), $\hat{u}$ can be expressed as

$$\hat{u} = \arg \min_u \min_{\boldsymbol{\theta}} R_{-k}([\boldsymbol{\theta}; u])$$

where

$$\begin{aligned} R_{-k}([\boldsymbol{\theta}; u]) &= \frac{1}{n} \sum_{i=1}^n (\boldsymbol{\theta}^\top \mathbf{r}_i + u e_i - y_i)^2 + \|\mathbf{J}\boldsymbol{\theta} + u \mathbf{h}\|_{\ell_2}^2 + 2g\left((\boldsymbol{\theta}^\top \mathbf{r}_i + u e_i - y_i)^2 \|\mathbf{J}\boldsymbol{\theta} + u \mathbf{h}\|_{\ell_2}^2\right) \\ (D.69) \quad &+ \lambda \|\boldsymbol{\theta}\|_{\ell_2}^2 + \lambda u^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} + \mathbf{w}_{N+1} u - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \lambda_s \frac{\sqrt{\log(d)}}{d} (\mathbf{1}^\top \boldsymbol{\theta} + u)^2, \end{aligned}$$

Here $\mathbf{h} \in \mathbb{R}^{N+1}$ is the last column of the $(N+1)$ -dimensional matrix $\mathbf{J}$. Namely,

$$(D.70) \quad \mathbf{h} := \begin{bmatrix} (\mathbf{W} \mathbf{w}_{N+1}) \odot \left(\frac{\pi - \cos^{-1}(\mathbf{W} \mathbf{w}_{N+1})}{2\pi} \right) \\ \frac{1}{2} \end{bmatrix}.$$

Let $f(u)$ denote the objective function of u above, i.e., $f(u) = \min_{\boldsymbol{\theta}} R_{-k}([\boldsymbol{\theta}; u])$. We also let $\boldsymbol{\theta}_*$ be the minimizing $\boldsymbol{\theta}$ in this objective if we set $u = 0$, i.e., $\boldsymbol{\theta}_* = \min_{\boldsymbol{\theta}} R_{-k}([\boldsymbol{\theta}; 0])$.

Following a similar argument in the proof of Proposition 6.2, we can obtain a lower bound on $f(u)$ by considering a second-order Taylor expansion of $f(u)$ around $[\boldsymbol{\theta}_*, 0]$ and using the strong-convexity of the loss function to arrive at the following upper bound on $\hat{u}$:

$$(D.71) \quad |\hat{u}| \leq \frac{1}{\lambda + \lambda_w} \left| \nabla_u R_{-k}([\boldsymbol{\theta}; u])|_{[\boldsymbol{\theta}_*, 0]} \right|.$$

Calculating $\nabla_u R_{-k}([\boldsymbol{\theta}; u])|_{[\boldsymbol{\theta}_*, 0]}$ we have

$$\begin{aligned} \nabla_u R_{-k}([\boldsymbol{\theta}; u])|_{[\boldsymbol{\theta}_*, 0]} &= \frac{1}{n} \sum_{i=1}^n \left[2(\boldsymbol{\theta}_*^\top \mathbf{r}_i - y_i) e_i + 2\mathbf{h}^\top \mathbf{J} \boldsymbol{\theta}_* \right. \\ &\quad \left. + \frac{2}{\sqrt{(\boldsymbol{\theta}_*^\top \mathbf{r}_i - y_i)^2 \|\mathbf{J} \boldsymbol{\theta}_*\|_{\ell_2}^2 + \gamma}} \left((\boldsymbol{\theta}_*^\top \mathbf{r}_i - y_i) \|\mathbf{J} \boldsymbol{\theta}_*\|_{\ell_2}^2 e_i + (\boldsymbol{\theta}_*^\top \mathbf{r}_i - y_i)^2 \mathbf{h}^\top \mathbf{J} \boldsymbol{\theta}_* \right) \right] \\ (D.72) \quad &+ 2\lambda_w \mathbf{w}_{N+1}^\top \left(\frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta}_* - \boldsymbol{\beta} \right) + 2\lambda_s \frac{\sqrt{\log(d)}}{d} \mathbf{1}^\top \boldsymbol{\theta}_*. \end{aligned}$$

We treat each of these terms separately.

Define the event

$$\mathcal{E} := \left\{ \frac{1}{n} \sum_{i=1}^n y_i^2 < C, \frac{1}{\sqrt{d}} \|\mathbf{X}\| \leq C, \|\mathbf{W}\| \leq C \right\}.$$

Using the concentration bounds for the operator norm of Gaussian matrices and also the tail bound for chi-square random variables, we have that $\mathbb{P}(\mathcal{E}) \geq 1 - 3e^{-cn}$ for some constant $c > 0$.

Note that by optimality of $\boldsymbol{\theta}_*$, we have $R_{-k}([\boldsymbol{\theta}_*; 0]) \leq R_{-k}([\mathbf{0}; 0]) = \frac{1}{n} \sum_{i=1}^n y_i^2$ from which we get $\|\boldsymbol{\theta}_*\|_{\ell_2} \leq C$ and $\left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta}_* - \boldsymbol{\beta} \right\|_{\ell_2} \leq C$, on the event $\mathcal{E}$. Therefore,

$$(D.73) \quad 2\lambda_s \frac{\sqrt{\log(d)}}{d} |\mathbf{1}^\top \boldsymbol{\theta}_*| \leq 2\lambda_s \frac{\sqrt{\log(d)}}{d} \|\boldsymbol{\theta}_*\|_{\ell_2} \|\mathbf{1}\|_{\ell_2} \leq C \sqrt{\frac{\log(d)}{d}}.$$

Also note that $\mathbf{w}_{N+1}$ is drawn independently from $\mathbf{W}$, $\boldsymbol{\theta}_*$ and $\boldsymbol{\beta}$. Since $\mathbf{w}_{N+1} \sim \text{Unif}(\mathbb{S}^{d-1})$, given $\frac{1}{2}\mathbf{W}^\top \boldsymbol{\theta}_* - \boldsymbol{\beta}$ the conditional distribution of $\mathbf{w}_{N+1}^\top (\frac{1}{2}\mathbf{W}^\top \boldsymbol{\theta}_* - \boldsymbol{\beta})$ converges to $\mathcal{N}(0, \frac{1}{d} \|\frac{1}{2}\mathbf{W}^\top \boldsymbol{\theta}_* - \boldsymbol{\beta}\|_{\ell_2}^2)$ from which we obtain

$$(D.74) \quad \left| \mathbf{w}_{N+1}^\top \left(\frac{1}{2}\mathbf{W}^\top \boldsymbol{\theta}_* - \boldsymbol{\beta} \right) \right| \leq \sqrt{2c' \frac{\log(d)}{d}} \left\| \frac{1}{2}\mathbf{W}^\top \boldsymbol{\theta}_* - \boldsymbol{\beta} \right\|_{\ell_2} < C \sqrt{2c' \frac{\log(d)}{d}},$$

with probability at least $1 - d^{-c'}$.

We next focus on the terms in the right-hand side of (D.72), which involve e_i . This part can be written as $\frac{1}{\sqrt{n}} \sum_{i=1}^n m_i e_i$ with

$$m_i = \frac{1}{\sqrt{n}} \left[2(\boldsymbol{\theta}_*^\top \mathbf{r}_i - y_i) + \frac{2}{\sqrt{(\boldsymbol{\theta}_*^\top \mathbf{r}_i - y_i)^2 \|\mathbf{J}\boldsymbol{\theta}_*\|_{\ell_2}^2 + \gamma}} (\boldsymbol{\theta}_*^\top \mathbf{r}_i - y_i) \|\mathbf{J}\boldsymbol{\theta}_*\|_{\ell_2}^2 \right].$$

Note that

$$|m_i| \leq \frac{2}{\sqrt{n}} (|\boldsymbol{\theta}_*^\top \mathbf{r}_i - y_i| + \|\mathbf{J}\boldsymbol{\theta}_*\|_{\ell_2}).$$

By optimality of $\boldsymbol{\theta}_*$, on the event $\mathcal{E}$ we have

$$\|\mathbf{m}\|_{\ell_2}^2 \leq \mathbb{R}_{-k}([\boldsymbol{\theta}_*; 0]) \leq \mathbb{R}_{-k}([0; 0]) = \frac{1}{n} \sum_{i=1}^n y_i^2 < C.$$

Observe that $\mathbf{w}_{N+1}$ is independent from $\{m_i\}_{i \in [n]}$ (recall that $\boldsymbol{\theta}_*$ does not depend on $\mathbf{w}_{N+1}$ by its definition.) Following the same strategy in the proof of Proposition 6.2, we only consider the randomness in $\mathbf{w}_{N+1}$ and condition on everything else. Write $\frac{1}{\sqrt{n}} \sum_{i=1}^n m_i e_i$ as a function of $\mathbf{w}_{N+1}$ as follows:

$$(D.75) \quad V(\mathbf{w}_{N+1}) := \frac{1}{\sqrt{n}} \sum_{i=1}^{k-1} m_i \sigma(\mathbf{w}_{N+1}^\top \mathbf{x}_i) + \frac{1}{\sqrt{n}} \sum_{i=k+1}^n m_i (\mu_1 \mathbf{w}_{N+1}^\top \mathbf{x}_i + \mu_2 z_i).$$

Observe that the (conditional) expectation $\mathbb{E}[V(\mathbf{w}_{N+1}) | \mathbf{W}, \mathbf{X}] = 0$. In addition, $V(\cdot)$ is a Lipschitz function with Lipschitz factor at most $\frac{C}{\sqrt{d}} \|\mathbf{X}\mathbf{m}\|_{\ell_2}$. Therefore, using the concentration bound for Lipschitz function on unit sphere (see e.g. [92, Theorem 5.1.4]), we obtain

$$\begin{aligned} \mathbb{P}(|V(\mathbf{w}_{N+1})| \geq t) &\leq \mathbb{P}(|V(\mathbf{w}_{N+1})| \geq t; \mathcal{E}) + \mathbb{P}(\mathcal{E}^c) \\ &\leq 2e^{-c't^2} + 3e^{-cn}. \end{aligned}$$

Choosing $t = C\sqrt{\frac{\log(d)}{d}}$, we get

$$(D.76) \quad \mathbb{P}\left(|V(\mathbf{w}_{N+1})| \geq C\sqrt{\frac{\log(d)}{d}}\right) \leq 2d^{-c'C^2} + 3e^{-cn}.$$

The remaining terms in (D.72) can be rearranged and written as $A\mathbf{h}^\top \mathbf{J}\boldsymbol{\theta}_*$, with

$$A := 2 \left(1 + \frac{1}{n} \sum_{i=1}^n \frac{(\boldsymbol{\theta}_*^\top \mathbf{r}_i - y_i)^2}{\sqrt{(\boldsymbol{\theta}_*^\top \mathbf{r}_i - y_i)^2 \|\mathbf{J}\boldsymbol{\theta}_*\|_{\ell_2}^2 + \gamma}} \right).$$

We next bound $A > 0$:

$$\begin{aligned}
|A| &< 2 \left(1 + \frac{1}{n} \sum_{i=1}^n \frac{|\boldsymbol{\theta}_*^\top \mathbf{r}_i - y_i|}{\|\mathbf{J}\boldsymbol{\theta}_*\|_{\ell_2}} \right) \\
&= 2 \left(1 + \frac{1}{\|\mathbf{J}\boldsymbol{\theta}_*\|_{\ell_2}} \left(\frac{1}{n} \sum_{i=1}^n (\boldsymbol{\theta}_*^\top \mathbf{r}_i - y_i)^2 \right)^{1/2} \right) \\
&< 2 \left(1 + \frac{1}{\|\mathbf{J}\boldsymbol{\theta}_*\|_{\ell_2}} (R_{-k}([\boldsymbol{\theta}_*; 0]))^{1/2} \right) \\
&\leq 2 \left(1 + \frac{1}{\|\mathbf{J}\boldsymbol{\theta}_*\|_{\ell_2}} (R_{-k}([\mathbf{0}; 0]))^{1/2} \right) \\
&= 2 \left(1 + \frac{1}{\|\mathbf{J}\boldsymbol{\theta}_*\|_{\ell_2}} \left(\frac{1}{n} \sum_{i=1}^n y_i^2 \right)^{1/2} \right).
\end{aligned}$$

Using Lemma F.3, on the event $\mathcal{E}$, the right-hand side of the above equation is of order one ($A < C$, for some constant $C > 0$).

We next bound $\mathbf{h}^\top \mathbf{J}\boldsymbol{\theta}_*$. Using the relation $\frac{1}{2\pi}(\pi - \cos^{-1}(\rho)) = \frac{1}{4} + \frac{\rho}{2\pi} + O(\rho^3)$, we define $\tilde{\mathbf{h}}$ as follows

$$\tilde{\mathbf{h}} := \begin{bmatrix} \frac{1}{4} \mathbf{W} \mathbf{w}_{n+1}^\top \\ \frac{1}{2} \end{bmatrix}.$$

Recalling $\mathbf{h}$ given by (D.70), we have $\|\mathbf{h} - \tilde{\mathbf{h}}\|_{\ell_2} = O(1/\sqrt{d})$. On the event $\mathcal{E}$, we have $\|\boldsymbol{\theta}_*\|_{\ell_2} = O(1)$. Also by invoking Lemma F.3, on event $\mathcal{E}$, we have $\|\mathbf{J}\| = O(1)$. Hence,

$$(D.77) \quad |\mathbf{h}^\top \mathbf{J}\boldsymbol{\theta}_* - \tilde{\mathbf{h}}^\top \mathbf{J}\boldsymbol{\theta}_*| \leq O(1/\sqrt{d}).$$

We henceforth focus on bounding $\tilde{\mathbf{h}}^\top \mathbf{J}\boldsymbol{\theta}_*$. Recall that $\mathbf{w}_{N+1}$ is independent of $\mathbf{W}$ and $\boldsymbol{\theta}_*$. Viewing $\tilde{\mathbf{h}}^\top \mathbf{J}\boldsymbol{\theta}_*$ as a function of $\mathbf{w}_{N+1}$, it has zero expectation (w.r.t $\mathbf{w}_{N+1}$ conditioned on $\boldsymbol{\theta}_*$ and $\mathbf{W}$). In addition, it is a Lipschitz function with Lipschitz factor at most $\frac{1}{4} \|\mathbf{W}^\top \mathbf{J}\boldsymbol{\theta}_*\|_{\ell_2}$, which is $O(1)$ on the event $\mathcal{E}$. Next, by employing the concentration bound for Lipschitz functions on unit sphere (see e.g. [92, Theorem 5.1.4]), we obtain

$$\begin{aligned}
\mathbb{P}(|\tilde{\mathbf{h}}^\top \mathbf{J}\boldsymbol{\theta}_*| \geq t) &\leq \mathbb{P}(|\tilde{\mathbf{h}}^\top \mathbf{J}\boldsymbol{\theta}_*| \geq t; \mathcal{E}) + \mathbb{P}(\mathcal{E}^c) \\
&\leq 2e^{-c't^2} + 3e^{-cn}.
\end{aligned}$$

Choosing $t = C\sqrt{\frac{\log(d)}{d}}$, and invoking (D.77) we get

$$(D.78) \quad \mathbb{P}\left(A |\mathbf{h}^\top \mathbf{J}\boldsymbol{\theta}_*| \geq C\sqrt{\frac{\log(d)}{d}}\right) \leq 2d^{-c'C^2} + 3e^{-cn}.$$

Combining the bounds (D.73) to (D.78) into (D.71) we get

$$|\hat{u}| \leq C' \frac{\log(d)}{\sqrt{d}},$$

with probability at least $4d^{-c'C^2} + 3e^{-cn} + d^{-c'}$.

The result follows by union bounding over the N coordinates of $\hat{\boldsymbol{\theta}}$, along with the assumption that N, n, d grow at the same order. (Note that the event $\mathcal{E}$ is common across all these bounds and so we count its complement probability once.) $\blacksquare$

D.2.5. Bounding the Difference of Φ_k, Φ_{k-1} with $\Psi_k(\mathbf{r})$

Lemma D.8 *We have*

$$(D.79) \quad \max \left\{ \mathbb{E}[(\Psi_k(\sigma(\mathbf{W}\mathbf{x}_k)) - \Phi_{-k})^2], (\Psi_k(\mathbf{f}_k) - \Phi_{-k})^2 \right\} \leq \frac{\text{polylog}(d)}{d^2},$$

and

$$(D.80) \quad \max \left\{ \mathbb{E}[(\Psi_k(\sigma(\mathbf{W}\mathbf{x}_k)) - \Phi_{k-1})^2], (\Psi_k(\mathbf{f}_k) - \Phi_k)^2 \right\} \leq \frac{\text{polylog}(d)}{d^3}$$

Proof To prove (D.79), we note from (D.13) that

$$\begin{aligned} \Psi_k(\mathbf{r}) - \Phi_{-k} &= \min_{\boldsymbol{\theta}} S_k(\boldsymbol{\theta}, \mathbf{r}) \leq S_k(\boldsymbol{\theta}_{-k}^*, \mathbf{r}) \\ &= \frac{1}{n} \ell(\boldsymbol{\theta}_{-k}^*; \mathbf{r}, y_k) \\ &\leq \frac{C}{n} (1 + |y_k| + \|\mathbf{J}\boldsymbol{\theta}_{-k}^*\|_{\ell_2} + \mathbf{r}^\top \boldsymbol{\theta}_{-k}^*)^2, \end{aligned}$$

where C is an absolute constant. Consequently, by using Lemma D.1 and the fact that $\|\mathbf{J}\|$ is bounded, we obtain (D.79).

To prove (D.80), we adapt the proof of Lemma 1 in [39] to our setting. In the following we use

$$q(\boldsymbol{\theta}) = \lambda \|\boldsymbol{\theta}\|_{\ell_2}^2 + \lambda_w \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \lambda_s \frac{\log(d)}{d} (\mathbf{1}^\top \boldsymbol{\theta})^2$$

We start with writing the Taylor expansion of $R_k(\boldsymbol{\theta}, \mathbf{r})$, defined in (D.9), around the point $\boldsymbol{\theta}_{-k}^*$. Note that $\boldsymbol{\theta}_{-k}^*$ is the minimizer of $R_{-k}(\boldsymbol{\theta})$, and hence

$$\begin{aligned} R_k(\boldsymbol{\theta}, \mathbf{r}) &= R_{-k}(\boldsymbol{\theta}_{-k}^*) + \frac{1}{n} \ell(\boldsymbol{\theta}; \mathbf{r}, y_k) + \frac{1}{2n} \sum_{t \neq k} (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)^\top \nabla^2 \ell(\boldsymbol{\theta}'; \mathbf{r}_t, y_t) (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*) \\ &\quad + \frac{1}{2} (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)^\top \nabla^2 q(\boldsymbol{\theta}') (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*) \end{aligned}$$

where $\boldsymbol{\theta}'$ can be written as

$$(D.81) \quad \boldsymbol{\theta}' = \omega \boldsymbol{\theta}_{-k}^* + (1 - \omega) \boldsymbol{\theta},$$

for some $\omega \in [0, 1]$. As a result, we can write using the definition (D.13)

$$\begin{aligned} (D.82) \quad R_k(\boldsymbol{\theta}, \mathbf{r}) - S_k(\boldsymbol{\theta}, \mathbf{r}) &= \frac{1}{2} (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)^\top \left[\frac{1}{n} \sum_{t \neq k} \nabla^2 \ell(\boldsymbol{\theta}'; \mathbf{r}_t, y_t) - \nabla^2 \ell(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t) \right] (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*), \end{aligned}$$

where we have noted that, since q is a quadratic function, we have $\nabla^2 q(\boldsymbol{\theta}_{-k}^*) = \nabla^2 q(\boldsymbol{\theta}')$.

Let us now consider the sum involving the terms of the form

$$(D.83) \quad (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)^\top (\nabla^2 \ell(\boldsymbol{\theta}'; \mathbf{r}_t, y_t) - \nabla^2 \ell(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t)) (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*).$$

We can now use the expansion in (D.21) to bound the above term. A straight-forward calculation, similar to what was done in the proof of Lemma D.4, shows that the above term involves several terms, among which the dominant term has the following form:

$$(D.84) \quad \alpha_t (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)^\top (\mathbf{r}_t \mathbf{r}_t^\top (\mathbf{r}_t^\top (\boldsymbol{\theta}' - \boldsymbol{\theta}_{-k}^*))) (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*) = \alpha_t (\mathbf{r}_t^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*))^2 (\mathbf{r}_t^\top (\boldsymbol{\theta}' - \boldsymbol{\theta}_{-k}^*)) = \alpha_t \omega (\mathbf{r}_t^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*))^3,$$

where α_t satisfies (noting that the derivatives of g are uniformly bounded):

$$(D.85) \quad \mathbb{P}(|\alpha_t| \geq v) \leq c \exp(-v^{c'}/c),$$

for absolute constants $c, c' > 0$. A straight-forward calculation (similar to what is done in the proof of Lemma D.4) shows that all the other terms in the expansion of (D.83) are in absolute value less than the term given in (D.84). As a result, one can write

$$\left| \frac{1}{2}(\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)^\top \left[\frac{1}{n} \sum_{t \neq k} \nabla^2 \ell(\boldsymbol{\theta}'_t; \mathbf{r}_t, y_t) - \nabla^2 \ell(\boldsymbol{\theta}_{-k}^*; \mathbf{r}_t, y_t) \right] (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*) \right| \leq \frac{1}{n} \sum_{t \neq k} |\alpha_t| |\mathbf{r}_t^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)|^3,$$

Using the above bound, we can now bound (D.82) as

$$(D.86) \quad |R_k(\boldsymbol{\theta}, \mathbf{r}) - S_k(\boldsymbol{\theta}, \mathbf{r})| \leq \frac{1}{n} \sum_{t \neq k} |\alpha_t| |\mathbf{r}_t^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_{-k}^*)|^3.$$

The rest of the proof follows almost line-by-line according to the proof of Lemma 1 in [39]. Let $\mathcal{B} = \{\boldsymbol{\theta}_k^*(\mathbf{r})\} \cup \{\tilde{\boldsymbol{\theta}}_k(\mathbf{r})\}$. By using the definitions (D.12) and (D.16), we have

$$\begin{aligned} |\Phi_k(\mathbf{r}) - \Psi_k(\mathbf{r})| &= \left| \min_{\boldsymbol{\theta} \in \mathcal{B}} R_k(\boldsymbol{\theta}, \mathbf{r}) - \min_{\boldsymbol{\theta} \in \mathcal{B}} S_k(\boldsymbol{\theta}, \mathbf{r}) \right| \\ &\leq \max_{\boldsymbol{\theta} \in \mathcal{B}} |R_k(\boldsymbol{\theta}, \mathbf{r}) - S_k(\boldsymbol{\theta}, \mathbf{r})| \end{aligned}$$

We thus obtain using (D.86) that

$$|\Phi_k(\mathbf{r}) - \Psi_k(\mathbf{r})| \leq C \frac{1}{n} \sum_{t \neq k} |\alpha_t| \left(|\mathbf{r}_t^\top (\boldsymbol{\theta}_k^*(\mathbf{r}) - \tilde{\boldsymbol{\theta}}_k(\mathbf{r}))|^3 + |\mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}}_k(\mathbf{r}) - \boldsymbol{\theta}_{-k}^*)|^3 \right)$$

Let us now bound each of the terms above. We have

$$\frac{1}{n} \sum_{t \neq k} |\alpha_t| |\mathbf{r}_t^\top (\boldsymbol{\theta}_k^*(\mathbf{r}) - \tilde{\boldsymbol{\theta}}_k(\mathbf{r}))|^3 \leq \|\boldsymbol{\theta}_k^*(\mathbf{r}) - \tilde{\boldsymbol{\theta}}_k(\mathbf{r})\|_{\ell_2}^3 \frac{1}{n} \sum_{t \neq k} |\alpha_t| \|\mathbf{r}_t\|_{\ell_2}^3.$$

Now, from Lemma D.3, up to negligible $O(\frac{\text{polylog}(n)}{n})$ terms, we have

$$\mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}}_k(\mathbf{r}) - \boldsymbol{\theta}_{-k}^*) = \frac{1}{n} \beta_1 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{r} + \frac{1}{n} \beta_2 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{p}.$$

Now, using the above relations, and the inequality $(\sum_{i=1}^n |a_i|)^2 \leq n \sum_{i=1}^n a_i^2$, as well as the Holder's inequality, we can write

$$\begin{aligned} |\Phi_k(\mathbf{r}) - \Psi_k(\mathbf{r})|^2 &\leq C'' \|\boldsymbol{\theta}_k^*(\mathbf{r}) - \tilde{\boldsymbol{\theta}}_k(\mathbf{r})\|_{\ell_2}^6 \left(\frac{1}{n} \sum_{t \neq k} |\alpha_t|^2 \|\mathbf{r}_t\|_{\ell_2}^6 \right) \\ &\quad + \frac{C''}{n} \sum_{t \neq k} |\alpha_t|^2 \left(\left| \beta_1 \frac{\mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{r}}{n} \right|^6 + \left| \beta_2 \frac{\mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{p}}{n} \right|^6 \right) \\ &\quad + O\left(\frac{\text{polylog}(n)}{n^4}\right), \end{aligned}$$

where $C'' > 0$ is an absolute constant. Now, from Lemma D.4, parts (c) and (e) of Lemma D.5, and (D.34), and (D.85), we obtain for any integer $D > 0$ that

$$\begin{aligned} \mathbb{E} \left[\|\boldsymbol{\theta}_k^*(\mathbf{r}) - \tilde{\boldsymbol{\theta}}_k(\mathbf{r})\|_{\ell_2}^{2D} \right] &\leq C_D \frac{\text{polylog}(n)}{n^{2D}}, \\ \mathbb{P} \left(\|\mathbf{r}_t\|_{\ell_2}^{2D} \geq n^D (\log(n))^b \right) &\leq c \exp(-(\log(d))^2/c), \\ \mathbb{P} \left(|\alpha_t|^{2D} \geq (\log(n))^b \right) &\leq c \exp(-(\log(d))^2/c), \\ \mathbb{P} \left(\max\{|\beta_1|^{2D}, |\beta_2|^{2D}\} \geq (\log(n))^b \right) &\leq c \exp(-(\log(d))^2/c), \end{aligned}$$

for suitably chosen absolute constants $b, c, C_D > 0$. Finally, since the matrix $\mathbf{H}_{-k}^{-1}$ has bounded norm, and $\mathbf{r}_t$ and $\mathbf{r}$ are independent sub-gaussian random vectors, we obtain

$$\mathbb{E} \left[\left| \frac{\mathbf{r}_t \mathbf{H}_{-k}^{-1} \mathbf{r}}{n} \right|^{2D} \right] \leq C_D \frac{\text{polylog}(n)}{n^D}$$

and $\mathbb{E} \left[\left| \frac{\mathbf{r}_t \mathbf{H}_{-k}^{-1} \mathbf{p}}{n} \right|^{2D} \right] \leq \mathbb{E} \left[\left| \frac{\|\mathbf{r}_t\|_{\ell_2} \|\mathbf{H}_{-k}^{-1}\|_{\ell_2} \|\mathbf{p}\|_{\ell_2}}{n} \right|^{2D} \right] \leq C_D \frac{\text{polylog}(n)}{n^D}.$

By using the above relations, the following result now follows in a straight-forward manner using the Holder's inequality:

$$\mathbb{E} \left[|\Phi_k(\mathbf{r}) - \Psi_k(\mathbf{r})|^2 \right] \leq \frac{\text{polylog}(n)}{n^3}.$$

And the result follows since d and n grow in proportion to each other. $\blacksquare$

D.2.6. Putting things together To prove Theorem 6.9, we consider any test function $\phi: \mathbb{R} \rightarrow \mathbb{R}$ which is uniformly bounded in terms of its value as well as its first and second derivatives. We will show that

$$(D.87) \quad |\mathbb{E}[\varphi(\Phi_A)] - \mathbb{E}[\varphi(\Phi_B)]| = \frac{\text{polylog}(d)}{d^{\frac{1}{2}}} + o_d(1).$$

Using this result, one immediately obtains the theorem (see Sections 2.3 and 2.4 of [39]). As a result, in the rest of this section we focus on proving the above relation for any test function φ . In order to prove this result, we use the so-called Lindeberg's method: We consider the quantities Φ_k defined in (D.7), and show for any $k \in [n]$ that

$$(D.88) \quad |\mathbb{E}[\varphi(\Phi_k)] - \mathbb{E}[\varphi(\Phi_{k-1})]| = \frac{\text{polylog}(d)}{d^{\frac{3}{2}}} + \frac{o_d(1)}{d}.$$

The above bound immediately results in (D.87) via a telescopic sum over k . It thus remains to prove (D.88).

Using the Taylor expansion, we can write

$$\varphi(\Phi_k) = \varphi(\Phi_{-k}) + \varphi'(\Phi_{-k})(\Phi_k - \Phi_{-k}) + \frac{1}{2}\varphi''(\alpha)(\Phi_k - \Phi_{-k})^2,$$

where α is a number between Φ_{-k} and Φ_k . Using the above expansion, and a similar expansion for Φ_{k-1} , we obtain

$$(D.89) \quad |\mathbb{E}[\varphi(\Phi_k)] - \mathbb{E}[\varphi(\Phi_{k-1})]| \leq \|\varphi'\|_{\infty} |\mathbb{E}[\Phi_k - \Phi_{k-1}]| + \frac{1}{2}\|\varphi''\|_{\infty} ((\Phi_k - \Phi_{-k})^2 + (\Phi_{k-1} - \Phi_{-k})^2),$$

where $\|\varphi'\|_{\infty}$ and $\|\varphi''\|_{\infty}$ are the maximum (absolute) values of the first and second derivative of φ .

By using Lemma D.8 we obtain that

$$\begin{aligned} |\mathbb{E}[\Phi_k - \Phi_{k-1}]| &\leq |\mathbb{E}[\Psi_k(\mathbf{f}_k) - \Psi_k(\sigma(\mathbf{W}\mathbf{x}_k))]| + \mathbb{E}[|\Phi_k - \Psi_k(\mathbf{f}_k)|] + \mathbb{E}[|\Psi_k(\sigma(\mathbf{W}\mathbf{x}_k)) - \Phi_{k-1}|] \\ &\leq |\mathbb{E}[\Psi_k(\mathbf{f}_k) - \Psi_k(\sigma(\mathbf{W}\mathbf{x}_k))]| + \frac{\text{polylog}(d)}{d^{\frac{3}{2}}}, \end{aligned}$$

where the last step follows simply from (D.80). Also, from (D.79) and (D.80), we can conclude that

$$\mathbb{E}[(\Phi_k - \Phi_{-k})^2] \leq 2\mathbb{E}[(\Phi_k - \Psi_k(\mathbf{f}_k))^2] + 2\mathbb{E}[(\Phi_{-k} - \Psi_k(\mathbf{f}_k))^2] \leq \frac{\text{polylog}(d)}{d^2},$$

and similarly

$$\mathbb{E}[(\Phi_{k-1} - \Phi_{-k})^2] \leq \frac{\text{polylog}(d)}{d^2}.$$

Finally, the only term that is left to be analyzed is $|\mathbb{E}[\Psi_k(\mathbf{f}_k) - \Psi_k(\sigma(\mathbf{W}\mathbf{x}_k))]|$, for which we use Lemma D.3, D.7, as well as a CLT-type result from [30]. We note the following three facts:

- (i) The quantity $\mathbf{r}^\top \boldsymbol{\theta}_{-k}^*$ converges in distribution to a gaussian with the same mean and variance when $\mathbf{r}$ is generated according to the distributions in (D.18). This is due to the CLT-type theorem given in [30, Theorem 2]. More precisely, we have shown in Lemma D.7 that with probability $1 - cd^{-c}$ we have: $\|\boldsymbol{\theta}_{-k}^*\|_{\ell_\infty}$ is at most $C\sqrt{(\log(d))/d}$, where c, C are absolute constants. Also, according to part (e) of Lemma D.1, we have $|\mathbf{1}^\top \boldsymbol{\theta}_{-k}^*| \leq C'\sqrt{d/(\log(d))}$, where C' is an absolute constant.

Now, let us define $\boldsymbol{\theta}' = \boldsymbol{\theta}_{-k}^*/\sqrt{\log(d)}$. Note that, with high probability (as specified above), we have $\|\boldsymbol{\theta}'\|_{\ell_\infty} \leq C/\sqrt{d}$ and $|\mathbf{1}^\top \boldsymbol{\theta}'| \leq C'\sqrt{d}/(\log(d))$. According to [30, Theorem 2], for $\mathbf{a}$ and $\mathbf{b}$ generated according (D.18), and fixing $\boldsymbol{\theta}'$, the random variables $\mathbf{a}^\top \boldsymbol{\theta}'$ and $\mathbf{b}^\top \boldsymbol{\theta}'$ have the same mean and variance, and we have

$$d_{\text{MS}}(\mathbf{a}^\top \boldsymbol{\theta}', \mathbf{b}^\top \boldsymbol{\theta}') \leq C'' \|\boldsymbol{\theta}'\|_{\ell_\infty} \left(\frac{|\mathbf{1}^\top \boldsymbol{\theta}'|}{\sqrt{d}} + \frac{1}{\sqrt{d}} \right),$$

where $C'' > 0$ is an absolute constant, and d_{MS} is the so-called maximum-sliced distance, and $d_{\text{MS}}(\mathbf{a}^\top \boldsymbol{\theta}', \mathbf{b}^\top \boldsymbol{\theta}')$ defines the distance between the distributions of $\mathbf{a}^\top \boldsymbol{\theta}'$ and $\mathbf{b}^\top \boldsymbol{\theta}'$. As a result, since $\boldsymbol{\theta}_{-k}^* = \boldsymbol{\theta}' \times \sqrt{\log(d)}$, we obtain that

$$d_{\text{MS}}(\mathbf{a}^\top \boldsymbol{\theta}_{-k}^*, \mathbf{b}^\top \boldsymbol{\theta}_{-k}^*) \leq C'' \sqrt{\log(d)} \|\boldsymbol{\theta}'\|_{\ell_\infty} \left(\frac{|\mathbf{1}^\top \boldsymbol{\theta}'|}{\sqrt{d}} + \frac{1}{\sqrt{d}} \right) = O\left(\frac{1}{\sqrt{\log(d)}}\right).$$

- (ii) Consider the result of Lemma D.3. From Lemma D.1, the norm of the vector $\boldsymbol{\theta}_{-k}^*$ is bounded by an absolute constant with probability at least $1 - \exp(-cn)$. Hence, since the matrix $\mathbf{J}$ is also of bounded operator norm, then the norm of the vector $\mathbf{p}$ given in Lemma D.3 is bounded by an absolute constant. Also, the quantity $\|\mathbf{J}\boldsymbol{\theta}_{-k}^*\|_{\ell_2}$ is bounded by an absolute constant. Given fixed matrix $\mathbf{H}^{-1}$ with bounded norm, and a fixed vector $\mathbf{p}$ with bounded norm, the quantity $\frac{1}{n}\mathbf{r}^\top \mathbf{H}^{-1}\mathbf{p}$ is, with probability at least $1 - c\exp(-(\log(d))^2/c)$, of order $O(\text{polylog}(d)/d)$ according to Lemma D.9. Hence, in the formula (D.31), the overall contribution of the terms which include $\frac{1}{n}\mathbf{r}^\top \mathbf{H}_{-k}^{-1}\mathbf{p}$ or $\frac{1}{n}\mathbf{p}^\top \mathbf{H}_{-k}^{-1}\mathbf{r}$ is of order $O(\text{polylog}(d)/d^2)$. Therefore, neglecting these terms adds an additional error of at most $O(\text{polylog}(d)/d^2)$ in computing $\mathbb{E}[\Psi(\mathbf{r})] - \Phi_{-k}$. Consequently, from the result of Lemma D.3 we can write

(D.90)

$$\mathbb{E}[\Psi_k(\mathbf{r})] = \Phi_{-k} + \frac{1}{n} \mathbb{E} \left[\min_{\tau_1} \left\{ \frac{\mathbf{r}^\top \mathbf{H}_{-k}^{-1} \mathbf{r}}{2n} \left(\frac{\partial \tilde{\ell}(\tau_1, 0)}{\partial \tau_1} \right)^2 + \tilde{\ell}(\tau_1, 0) \right\} \right] + O\left(\frac{\text{polylog}(d)}{d^{\frac{3}{2}}}\right),$$

where $\tilde{\ell}(\tau_1, \tau_2)$ is given in (D.32).

- (iii) Given a matrix $\mathbf{H}$, the value $\frac{1}{n}\mathbf{r}^\top \mathbf{H}_{-k}^{-1}\mathbf{r}$ concentrates on the same quantity if $\mathbf{r}$ is generated from either of the distributions in (D.18). More precisely, from [39, Lemma 13] (or [56, Lemma 1]) we obtain

$$\mathbb{P} \left(\left| \frac{1}{n} \mathbf{r}^\top \mathbf{H}_{-k}^{-1} \mathbf{r} - \mathbb{E} \left[\frac{1}{n} \mathbf{r}^\top \mathbf{H}_{-k}^{-1} \mathbf{r} \right] \right| \geq c \frac{\log(d)}{\sqrt{d}} \right) \leq 1 - c \exp(-(\log(d))^2),$$

for $\mathbf{r}$ being generated according to either of the distributions in (D.18), and $c > 0$ being an absolute constant. Also, we have that

$$\left| \mathbb{E} \left[\frac{1}{n} \sigma(\mathbf{W}\mathbf{x})^\top \mathbf{H}_{-k}^{-1} \sigma(\mathbf{W}\mathbf{x}) \right] - \mathbb{E} \left[\frac{1}{n} \mathbf{f}^\top \mathbf{H}_{-k}^{-1} \mathbf{f} \right] \right| = \frac{1}{n} \text{Trace} \left(\mathbf{H}_{-k}^{-1} (\Sigma_s - \Sigma_f) \right),$$

where $\Sigma_2 = \mathbb{E}[\sigma(\mathbf{W}\mathbf{x})\sigma(\mathbf{W}\mathbf{x})^\top]$ and $\mathbb{E}[\mathbf{f}\mathbf{f}^\top]$, and the last inequality follows from Lemma D.10. As a result, we obtain

$$\begin{aligned} \mathbb{P} \left(\left| \frac{1}{n} \sigma(\mathbf{W}\mathbf{x}_k)^\top \mathbf{H}_{-k}^{-1} \sigma(\mathbf{W}\mathbf{x}_k) - \frac{1}{n} \mathbb{E}[\mathbf{f}^\top \mathbf{H}_{-k}^{-1} \mathbf{f}] \right| \geq c \frac{(\log(d))^{3/2}}{\sqrt{d}} \right) &\leq 1 - c \exp(-(\log(d))^2), \\ \mathbb{P} \left(\left| \frac{1}{n} \mathbf{f}^\top \mathbf{H}_{-k}^{-1} \mathbf{f} - \frac{1}{n} \mathbb{E}[\mathbf{f}^\top \mathbf{H}_{-k}^{-1} \mathbf{f}] \right| \geq c \frac{\log(d)}{\sqrt{d}} \right) &\leq 1 - c \exp(-(\log(d))^2), \end{aligned}$$

Let us now put all the above facts together to bound $|\mathbb{E}[\Psi_k(\mathbf{f}_k) - \Psi_k(\sigma(\mathbf{W}\mathbf{x}_k))]|$. Consider the function $\tilde{\ell}(\tau_1, \tau_2)$ given in (D.32). This function depends on $\mathbf{r}$ only through $\boldsymbol{\theta}_{-k}^* \mathbf{r}$. Using fact (i) above, we know that $\boldsymbol{\theta}_{-k}^* \mathbf{r}$ will asymptotically have the same (gaussian) distribution for both $\mathbf{r} \sim \mathbf{f}_k$ and $\mathbf{r} \sim \sigma(\mathbf{W}\mathbf{x}_k)$. Also, it is easy to conclude using part (c) of Lemma D.1 that all of the moments of random variable $\boldsymbol{\theta}_{-k}^* \mathbf{r}$ are bounded (i.e. $\mathbb{E}[|\boldsymbol{\theta}_{-k}^* \mathbf{r}|^D] \leq C_D$ for an absolute constant $C_D > 0$). Further, in fact (ii) we have argued that $\|J\boldsymbol{\theta}_{-k}^*\|_{\ell_2}$ is bounded with probability $1 - e \exp(-cn)$. Also, from fact (iii) above, we know that the term $\frac{1}{n} \mathbf{r}^\top \mathbf{H}_{-k}^{-1} \mathbf{r}$ concentrates sharply on the same value for $\mathbf{r}$ being either $\mathbf{f}_k$ or $\sigma(\mathbf{W}\mathbf{x}_k)$. Putting all these together, and using [7, Corollary of Theorem 25.12], we obtain that

$$\begin{aligned} &\mathbb{E}_{\mathbf{r}=\mathbf{f}_k} \left[\min_{\tau_1} \left\{ \frac{\mathbf{r}^\top \mathbf{H}_{-k}^{-1} \mathbf{r}}{2n} \left(\frac{\partial \tilde{\ell}(\tau_1, 0)}{\partial \tau_1} \right)^2 + \tilde{\ell}(\tau_1, 0) \right\} \right] \\ &\quad - \mathbb{E}_{\mathbf{r}=\sigma(\mathbf{W}\mathbf{x}_k)} \left[\min_{\tau_1} \left\{ \frac{\mathbf{r}^\top \mathbf{H}_{-k}^{-1} \mathbf{r}}{2n} \left(\frac{\partial \tilde{\ell}(\tau_1, 0)}{\partial \tau_1} \right)^2 + \tilde{\ell}(\tau_1, 0) \right\} \right] \\ &= o_d(1), \end{aligned}$$

and therefore from (D.90) we obtain

$$|\mathbb{E}[\Psi_k(\mathbf{f}_k) - \Psi_k(\sigma(\mathbf{W}\mathbf{x}_k))]| \leq \frac{o_d(1)}{d},$$

for an absolute constant $C > 0$, and hence we obtain (D.88).

D.2.7. Proofs of the Auxiliary Lemmas Here we provide the proofs of some of the auxiliary lemmas used in our analysis.

Proof of Lemma D.1. We will prove part (b) here, but part (a) will have the exact same proof. Let $\boldsymbol{\theta}^*$ be the minimizer of $R_k(\boldsymbol{\theta}, \mathbf{r})$. We can write

$$R_k(\boldsymbol{\theta}^*, \mathbf{r}) \leq R_k(\mathbf{0}, \mathbf{r}),$$

as $\boldsymbol{\theta}^*$ is the minimizer.

On the one hand we have

$$R_k(\mathbf{0}, \mathbf{r}) = \frac{1}{n} \sum_{i=1}^n y_i^2 + \|\boldsymbol{\beta}\|_{\ell_2}^2,$$

and thus for any $v \geq 0$:

$$(D.91) \quad \mathbb{P}\left(R_k(\mathbf{0}, \mathbf{r}) \geq v + \|\boldsymbol{\beta}\|_{\ell_2}^2 + 2\mathbb{E}[y_1^2]\right) \leq c_1 \exp(-nv^2/c_1),$$

for an absolute constant $c_1 > 0$.

On the other hand, since $R(\boldsymbol{\theta}, \mathbf{r})$ is λ -strongly convex, and $R(\boldsymbol{\theta}, \mathbf{r}) \geq 0$, we can write

$$\|\boldsymbol{\theta}^*\|_{\ell_2}^2 \leq \frac{1}{\lambda} R(\mathbf{0}, \mathbf{r}),$$

which together with (D.91) gives use the result.

To prove part (c), we note that $\mathbf{r}$ is generated independently from $\boldsymbol{\theta}_{-k}^*$. We can thus write:

$$\begin{aligned} \mathbb{P}\left(|\mathbf{r}^\top \boldsymbol{\theta}_{-k}^*| \geq v\right) &\leq \mathbb{P}\left(\|\boldsymbol{\theta}_{-k}^*\|_{\ell_2} \geq \sqrt{v}\right) + \mathbb{P}\left(|\mathbf{r}^\top \boldsymbol{\theta}_{-k}^*| \geq v \mid \|\boldsymbol{\theta}_{-k}^*\|_{\ell_2} < \sqrt{v}\right) \\ &\leq c_2 e^{-v/c_2} + \mathbb{P}\left(|\mathbf{r}^\top \boldsymbol{\theta}_{-k}^*| \geq v \mid \|\boldsymbol{\theta}_{-k}^*\|_{\ell_2} < \sqrt{v}\right) \\ &\leq c' e^{-v/c'}. \end{aligned}$$

where the second step follows from part (a) of the lemma (with c_2 chosen to be sufficiently large); and the last step follows from the independence of $\mathbf{r}$ and $\boldsymbol{\theta}_{-k}^*$ as well as Lemma D.9.

Also, the proof of part (d) follows simply because

$$\lambda_s \frac{d}{\log(d)} (\mathbf{1}^\top \boldsymbol{\theta}_{-k}^*)^2 \leq R_{-k}(\boldsymbol{\theta}_{-k}^*) \leq R_{-k}(\mathbf{0}) = \frac{1}{n} \sum_{i \neq k} y_i^2 + \|\boldsymbol{\beta}\|_{\ell_2}^2,$$

where $\mathbf{0}$ is the all-zero vector. By using a bound similar to (D.91) for $R_{-k}(\mathbf{0})$ we obtain the result.

Proof of Lemma D.5. Part (a) is exactly Lemma 12 in [39]. For part (b), consider the matrix $\mathbf{R}$ whose columns are $\mathbf{r}_t$'s, i.e. $\mathbf{R} = [\mathbf{r}_1 | \mathbf{r}_2 | \dots | \mathbf{r}_n]$. Since $\mathbf{r}_t$'s are zero-mean sub-gaussian vectors (see Lemma D.9), we know that its operator norm satisfies:

$$\mathbb{P}\left(\|\mathbf{R}\| \geq c_1 \sqrt{d} + v\right) \leq c \exp(-v^2/c).$$

Also, define the vector $\alpha = [\alpha_t]_{t \neq k}^\top$. Note that $\|\alpha\|_{\ell_2} \leq \sqrt{n}$. We have

$$\mathbb{P}\left(\left\|\frac{1}{n} \sum_{t \neq k} \alpha_t \mathbf{r}_t\right\|_{\ell_2} \geq v + c_1\right) = \mathbb{P}\left(\left\|\frac{1}{n} \mathbf{R} \alpha\right\|_{\ell_2} \geq v + c_1\right) = \mathbb{P}\left(\|\mathbf{R}\| \geq v\sqrt{n} + c_1\sqrt{n}\right).$$

Given the above relation, and the fact that d and n grow proportionally, the result of the second part of the lemma follows easily.

Part (c) follows from Lemma D.9. Also, part(d) follows from part (c) as well as Lemma D.2 (specifically (D.33)) and the fact that the operator norm of the matrix $\mathbf{H}_{-k}^{-1}$ is upper-bounded by $2/\lambda$.

Part (e) follows from the fact that $\mathbf{r}$ is a random sub-gaussian vector (see Lemma D.9). We refer to [91] for bounds on the ℓ_2 norm of random sub-gaussian vectors.

To prove part (e), we use (D.33) to write

$$\begin{aligned} &\mathbb{P}\left(|\mathbf{r}_t^\top (\tilde{\boldsymbol{\theta}}(\mathbf{r}) - \boldsymbol{\theta}_{-k}^*)| \geq \frac{v}{\sqrt{n}}\right) \leq \mathbb{P}\left(\frac{1}{n} (|\beta_1 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{r}| + |\beta_2 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{p}|) + |\mathbf{r}_t^\top \mathbf{e}| \geq \frac{v}{\sqrt{n}}\right), \\ (D.92) \quad &\leq \mathbb{P}\left(\frac{1}{n} |\beta_1 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{r}| \geq \frac{v}{3\sqrt{n}}\right) + \mathbb{P}\left(\frac{1}{n} |\beta_2 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{p}| \geq \frac{v}{3\sqrt{n}}\right) + \mathbb{P}\left(|\mathbf{r}_t^\top \mathbf{e}| \geq \frac{v}{3\sqrt{n}}\right) \end{aligned}$$

We will now bound each of the terms above. For the first term we have

$$\begin{aligned}
& \mathbb{P}\left(\frac{1}{n}|\beta_1 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{r}| \geq \frac{v}{3\sqrt{n}}\right) \\
& \leq \mathbb{P}(\|\mathbf{r}_t\|_{\ell_2} \geq \sqrt{vn}) + \mathbb{P}\left(\frac{1}{n}|\beta_1 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{r}| \geq \frac{v}{3\sqrt{n}}, \|\mathbf{r}_t\|_{\ell_2} < \sqrt{vn}\right) \\
& \leq c_1 \exp(-v^{c_2}/c_1) + \mathbb{P}\left(\frac{1}{n}|\beta_1 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{r}| \geq \frac{v}{3\sqrt{n}}, \|\mathbf{r}_t\|_{\ell_2} < \sqrt{vn}\right),
\end{aligned}$$

where the last step follows from part (e) and appropriately selecting $c_1, c_2 > 0$. Now, note that the vector $\mathbf{r}$ is generated independently from $\mathbf{r}_t$ and $\mathbf{H}_{-k}$. As a result, to bound the second term in the RHS of the above relation, we notice that, assuming $\|\mathbf{r}_t\|_{\ell_2} < \sqrt{vn}$, we have $\|\beta_1 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1}\|_{\ell_2} \leq C|\beta_1|\sqrt{vn}$. Hence, we can write

$$\begin{aligned}
& \mathbb{P}\left(\frac{1}{n}|\beta_1 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{r}| \geq \frac{v}{3\sqrt{n}}, \|\mathbf{r}_t\|_{\ell_2} < \sqrt{vn}\right) \\
& \leq \mathbb{P}(|\beta_1| \geq v^{1/4}) + \mathbb{P}\left(\frac{1}{n}|\beta_1 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{r}| \geq \frac{v}{3\sqrt{n}}, \|\mathbf{r}_t\|_{\ell_2} < \sqrt{vn}, |\beta_1| < v^{1/4}\right) \\
& \leq c_3 \exp(-v^{c_4}/c_3) + \mathbb{P}\left(\frac{1}{n}|\beta_1 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{r}| \geq \frac{v}{3\sqrt{n}}, \|\mathbf{r}_t\|_{\ell_2} < \sqrt{vn}, |\beta_1| < v^{1/4}\right) \\
& \leq c_3 \exp(-v^{c_4}/c_3) + c_5 \exp(-v^{c_6}/c_5),
\end{aligned}$$

where the second inequality follows from (D.34) and by suitably choosing $c_3, c_4 > 0$. The third inequality follows from sub-gaussianity of $\mathbf{r}$, see Lemma D.9, and the fact that, given $|\beta_1| < v^{1/4}$ and $\|\mathbf{r}_t\|_{\ell_2} \leq \sqrt{vn}$, the random vector $\mathbf{u} = \beta_1 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1}$ satisfies $\|\mathbf{u}\|_{\ell_2} \leq C v^{3/4} \sqrt{n}$ and is independently generated from $\mathbf{r}$. Hence, we can bound the second term in the RHS of the above relation by using Lemma D.9 and appropriate choices of $c_5, c_6 > 0$.

The second and third terms in (D.92) can be bounded similarly as the first term but in an easier manner. The second term follows by writing $|\mathbf{r}_t^\top \mathbf{p}| \leq \|\mathbf{r}_t\|_{\ell_2} \|\mathbf{p}\|_{\ell_2}$, and noticing that $\|\mathbf{p}\|_{\ell_2}$ is upper-bounded by a constant since the norm of $\mathbf{J}$ is bounded and the norm of $\boldsymbol{\theta}_{-k}^*$ is bounded (see Lemma D.3 and Lemma D.1). Hence, using a similar (but simpler) argument as above, we can write

$$\mathbb{P}\left(\frac{1}{n}|\beta_2 \mathbf{r}_t^\top \mathbf{H}_{-k}^{-1} \mathbf{p}| \geq \frac{v}{3\sqrt{n}}\right) \leq c_7 \exp(-v^{c_8}/c_7),$$

for absolute constants $c_7, c_8 > 8$.

Finally, the third term in the RHS of (D.92) can be bounded by writing $|\mathbf{r}_t^\top \mathbf{e}| \leq \|\mathbf{r}_t\|_{\ell_2} \|\mathbf{e}\|_{\ell_2}$, and noticing that $\|\mathbf{e}\|_{\ell_2}$ is small according to (D.34). And, using similar steps as above, we reach to a similar upper bound.

Part (g) follows from Lemma D.1 and Lemma D.6.

Lemma D.9 [Analogous to Lemma 8 in [39]] Assume that $\mathbf{a} = \sigma(\mathbf{W}\mathbf{x})$ and $\mathbf{b} = \mu_1 \mathbf{W}\mathbf{x} + \mu_2 \mathbf{u}$, where $\mathbf{x}$ and $\mathbf{u}$ are generated independently from the normal distribution. Also, let $\boldsymbol{\Sigma} = \mathbb{E}[\mathbf{b}\mathbf{b}^\top]$, i.e. $\boldsymbol{\Sigma} = \mu_1^2 \mathbf{W}\mathbf{W}^\top + \mu_2^2 \mathbf{I}$. Then, there exists an absolute constant $c > 0$ such that:

$$(D.93) \quad \mathbb{P}(|\mathbf{a}^\top \boldsymbol{\beta}| \geq v) \leq 2 \exp\left(-\frac{v^2}{c \|\boldsymbol{\beta}\|_{\ell_2}^2 \|\mathbf{W}\|^2}\right),$$

and

$$(D.94) \quad \mathbb{P}(|\mathbf{b}^\top \boldsymbol{\beta}| \geq v) \leq 2 \exp\left(-\frac{v^2}{c \|\boldsymbol{\beta}\|_{\ell_2}^2 \|\boldsymbol{\Sigma}\|}\right),$$

for a fixed vector $\boldsymbol{\beta} \in \mathbb{R}^d$ and any $v \geq 0$. Here, $\|\mathbf{W}\|$ (resp. $\|\boldsymbol{\Sigma}\|$) denotes the operator norm of $\mathbf{W}$ (resp. $\boldsymbol{\Sigma}$).

Proof We will be using the following well-known relation: For a L -Lipschitz continuous function f and $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ we have

$$(D.95) \quad \mathbb{P}(|f(\mathbf{x}) - \mathbb{E}[f(\mathbf{x})]| \geq v) \leq 2 \exp\left(-\frac{v^2}{4L^2}\right).$$

Now, note that since σ is the shifted Relu function, it is easy to see that the function $f(\mathbf{x}) = \boldsymbol{\beta}^\top \sigma(\mathbf{W}\mathbf{x})$ is $\|\boldsymbol{\beta}\|_{\ell_2} \|\mathbf{W}\|$ -Lipschitz continuous. Therefore, we obtain the result using the relation (D.95) and the fact that $\mathbf{x}$ is distributed according to the normal distribution.

The proof of the second part can similarly be done by noting that $\mathbf{b} = \boldsymbol{\Sigma}^{1/2} \tilde{\mathbf{b}}$ where $\tilde{\mathbf{b}}$ is distributed according to the standard normal distribution. $\blacksquare$

Lemma D.10 Let $\mathbf{H} \in \mathbb{R}^{N \times N}$ be such that $\|\mathbf{H}\| \leq C$ for an absolute constant $C > 0$. Let $\boldsymbol{\Sigma}_s = \mathbb{E}[\sigma(\mathbf{W}\mathbf{x})\sigma(\mathbf{W}\mathbf{x})^\top]$ and $\boldsymbol{\Sigma}_f = \mathbb{E}[\mathbf{f}\mathbf{f}^\top]$. We have

$$(D.96) \quad \left| \frac{1}{n} \text{Trace} \{ \mathbf{H}(\boldsymbol{\Sigma}_s - \boldsymbol{\Sigma}_f) \} \right| \leq \frac{c(\log(d))^{3/2}}{\sqrt{d}}.$$

with probability at least $1 - c \exp(-(\log(d))^2/c)$ where $c > 0$ is an absolute constant.

Proof We first bound each element of the matrix $\boldsymbol{\Sigma}_s - \boldsymbol{\Sigma}_f$. Recall that the k -th element of the vector $\mathbf{f}$ is distributed according to

$$\mu_1 \mathbf{w}_k^\top \mathbf{x} + \mu_2 u_k,$$

where u_k is independently generated from $\mathcal{N}(0, 1)$ and $\mu_1 = \frac{1}{2}, \mu_2 = \sqrt{\frac{1}{4} - \frac{1}{2\pi}}$. As a result, the (ℓ, k) -th element of the matrix $\boldsymbol{\Sigma}_f$ is

$$(D.97) \quad \mathbb{E}[(\mu_1 \mathbf{w}_k^\top \mathbf{x} + \mu_2 u_k)(\mu_1 \mathbf{w}_\ell^\top \mathbf{x} + \mu_2 u_\ell)] = \mu_1^2 \mathbf{w}_k^\top \mathbf{w}_\ell + \mu_2^2 \mathbb{I}\{k = \ell\}.$$

Note that $\langle \mathbf{w}_\ell, \mathbf{x} \rangle$ and $\langle \mathbf{w}_k, \mathbf{x} \rangle$ are jointly Gaussian with

$$\mathbb{E}[(\mathbf{w}_\ell^\top \mathbf{x})^2] = \mathbb{E}[(\mathbf{w}_k^\top \mathbf{x})^2] = 1, \quad \mathbb{E}[(\mathbf{w}_\ell^\top \mathbf{x})(\mathbf{w}_k^\top \mathbf{x})] = \mathbf{w}_k^\top \mathbf{w}_\ell.$$

Therefore, we have (see e.g., [15, Table 1])

$$(D.98) \quad \begin{aligned} \mathbb{E}[\sigma(\mathbf{w}_\ell^\top \mathbf{x})\sigma(\mathbf{w}_k^\top \mathbf{x})] &= \frac{\sqrt{1 - (\mathbf{w}_k^\top \mathbf{w}_\ell)^2} + (\pi - \cos^{-1}(\mathbf{w}_k^\top \mathbf{w}_\ell))(\mathbf{w}_k^\top \mathbf{w}_\ell)}{2\pi} - \frac{1}{2\pi} \\ &= \frac{1}{4} \mathbf{w}_\ell^\top \mathbf{w}_k + \left(\frac{1}{4} - \frac{1}{2\pi}\right) \mathbb{I}\{k = \ell\} + O((\mathbf{w}_\ell^\top \mathbf{w}_k)^3) \\ &= \mu_1^2 \mathbf{w}_\ell^\top \mathbf{w}_k + \mu_2^2 \mathbb{I}\{k = \ell\} + O\left(\left(\frac{\log(d)}{d}\right)^{\frac{3}{2}}\right), \end{aligned}$$

where the last step follows from the fact that with probability at least $1 - c \exp(-(\log(d))^2/c)$ we have for all k, ℓ , such that $k \neq \ell$, we have $|\mathbf{w}_\ell^\top \mathbf{w}_k| \leq \frac{\log(d)}{d}$. As a result, from (D.97) and (D.98) we obtain

$$|(\boldsymbol{\Sigma}_s - \boldsymbol{\Sigma}_f)_{k,\ell}| = O\left(\left(\frac{\log(d)}{d}\right)^{\frac{3}{2}}\right).$$

Hence, the ℓ_2 norm of each column of the matrix $\Sigma_s - \Sigma_f$ is of order $O((\log(d))^{3/2}/d)$. Now, since $\|\mathbf{H}\| \leq C$, then the ℓ_2 norm of each row of $\mathbf{H}$ is at most C . Thus, by a simple application of the Cauchy-Schwarz inequality we obtain the result of the lemma. $\blacksquare$

D.3. Proof of Proposition 6.10 Let us first show that $\widehat{\boldsymbol{\theta}}^*, \widehat{\boldsymbol{\theta}}_{\text{nl}}^*$ fall in $\mathcal{C}_{-\boldsymbol{\theta}}$ with high probability for any $\zeta > 0$. We prove the result for $\widehat{\boldsymbol{\theta}}^*$, and remark that the proof is exactly the same for $\widehat{\boldsymbol{\theta}}_{\text{nl}}^*$. Consider the objective

$$R(\boldsymbol{\theta}) := \frac{1}{2n} \sum_{i=1}^n (|y_i - \boldsymbol{\theta}^\top \sigma(\mathbf{W} \mathbf{x}_i)| + \varepsilon \|\mathbf{J} \boldsymbol{\theta}\|_{\ell_2})^2 + \frac{\zeta}{2} \boldsymbol{\theta}^\top \boldsymbol{\Omega} \boldsymbol{\theta}.$$

On the one hand we have

$$R(\mathbf{0}) = \frac{1}{n} \sum_{i=1}^n y_i^2$$

where $\mathbf{0}$ is the all-zero vector. Thus for any $v \geq 0$:

$$(D.99) \quad \mathbb{P}(R_k(\mathbf{0}) \geq v + 2\mathbb{E}[y_1^2]) \leq c_1 \exp(-nv^2/c_1),$$

for an absolute constant $c_1 > 0$.

On the other hand, since $R(\boldsymbol{\theta})$ is ζ -strongly convex, and $R(\boldsymbol{\theta}) \geq 0$, we can write

$$\|\widehat{\boldsymbol{\theta}}^*\|_{\ell_2}^2 \leq \frac{1}{\zeta} R(\mathbf{0}, \mathbf{r}),$$

which together with (D.99) proves that $\|\widehat{\boldsymbol{\theta}}^*\|_{\ell_2}$ is bounded above by a constant C with probability at least $1 - e \exp(-cn)$.

To bound $|\mathbf{1}^\top \widehat{\boldsymbol{\theta}}^*|$ we note that

$$\zeta \frac{d}{\log(d)} (\mathbf{1}^\top \widehat{\boldsymbol{\theta}}^*)^2 \leq R(\widehat{\boldsymbol{\theta}}^*) \leq R(\mathbf{0}) = \frac{1}{n} \sum_{i \neq k} y_i^2,$$

By using the bound to (D.99) we obtain with probability $1 - e \exp(-cn)$ that $|\mathbf{1}^\top \widehat{\boldsymbol{\theta}}^*| \leq C\sqrt{d/(\log(d))}$. Finally, the fact that $\|\widehat{\boldsymbol{\theta}}^*\|_{\ell_\infty}$ is with high probability of order $\sqrt{(\log(d))/d}$ follows in exactly the same manner as the proof of Lemma D.7 and hence we do not repeat the proof here.

We now show that $M(\widehat{\boldsymbol{\theta}}^*) - M(\widehat{\boldsymbol{\theta}}_{\text{nl}}^*) \rightarrow 0$ in probability. To do so, we use the argument given in [1, Theorem 4]. Assume that $M(\widehat{\boldsymbol{\theta}}^*)$ and $M(\widehat{\boldsymbol{\theta}}_{\text{nl}}^*)$ converge to different values, say M_A and M_B . Define $M = (M_A + M_B)/2$ and consider the following optimization problems

$$\begin{aligned} \bar{\Phi}_A &:= \min_{\boldsymbol{\theta}: M(\boldsymbol{\theta}) \leq M} \frac{1}{2n} \sum_{i=1}^n (|y_i - \boldsymbol{\theta}^\top \sigma(\mathbf{W} \mathbf{x}_i)| + \varepsilon \|\mathbf{J} \boldsymbol{\theta}\|_{\ell_2})^2 + \frac{\zeta}{2} \boldsymbol{\theta}^\top \boldsymbol{\Omega} \boldsymbol{\theta}, \\ \bar{\Phi}_B &:= \min_{\boldsymbol{\theta}: M(\boldsymbol{\theta}) \leq M} \frac{1}{2n} \sum_{i=1}^n (|y_i - \boldsymbol{\theta}^\top \mathbf{f}_i| + \varepsilon \|\mathbf{J} \boldsymbol{\theta}\|_{\ell_2})^2 + \frac{\zeta}{2} \boldsymbol{\theta}^\top \boldsymbol{\Omega} \boldsymbol{\theta}. \end{aligned}$$

Note that the values $\bar{\Phi}_A$ and $\bar{\Phi}_B$ must be different. Now, using the minimax theorem, and since the above objectives are ζ -strongly convex, we can write

$$\begin{aligned} \bar{\Phi}_A &= \sup_{\lambda > 0} -\lambda M + \min_{\boldsymbol{\theta}} \frac{1}{2n} \sum_{i=1}^n (|y_i - \boldsymbol{\theta}^\top \sigma(\mathbf{W} \mathbf{x}_i)| + \varepsilon \|\mathbf{J} \boldsymbol{\theta}\|_{\ell_2})^2 + \frac{\zeta}{2} \boldsymbol{\theta}^\top \boldsymbol{\Omega} \boldsymbol{\theta} + \lambda M(\boldsymbol{\theta}), \\ \bar{\Phi}_B &= \sup_{\lambda > 0} -\lambda M + \min_{\boldsymbol{\theta}} \frac{1}{2n} \sum_{i=1}^n (|y_i - \boldsymbol{\theta}^\top \mathbf{f}_i| + \varepsilon \|\mathbf{J} \boldsymbol{\theta}\|_{\ell_2})^2 + \frac{\zeta}{2} \boldsymbol{\theta}^\top \boldsymbol{\Omega} \boldsymbol{\theta} + \lambda M(\boldsymbol{\theta}). \end{aligned}$$

Now, from the result of Theorem 6.9 we know that for any $\lambda > 0$ the values inside the min converge to the same value. As a result, the quantities $\bar{\Phi}_A$ and $\bar{\Phi}_B$ should converge to the same value (according to [1, Lemma 1]) which is a contradiction with the claim that $M(\hat{\theta}^*)$ and $M(\hat{\theta}_{\text{nl}}^*)$ converge to different values (i.e. M_A and M_B , respectively). A similar argument can be applied to show that $\|\mathbf{J}\hat{\theta}^*\|_{\ell_2} - \|\mathbf{J}\hat{\theta}_{\text{nl}}^*\|_{\ell_2} \rightarrow 0$, in probability.

APPENDIX E: PROOFS OF STEP 4: ANALYSIS OF THE GAUSSIAN NOISY LINEAR MODEL VIA CONVEX GAUSSIAN MINIMAX FRAMEWORK

By the Gaussian equivalence property, we henceforth focus on optimization (6.23) and provide a precise characterization of $\overset{\circ}{\text{AR}}_{\text{nl}}(\hat{\theta}_{\text{nl}}^*)$.

Before proceeding, we will discuss another representation of the model using a few change of variables. Recall from (6.15) that $\mathbf{f} := \mu_0 \mathbf{1} + \mu_1 \mathbf{W} \mathbf{x} + \mu_2 \mathbf{u}$. For our activation function $\sigma(v) = v \mathbb{1}(v \geq 0) - 1/\sqrt{2\pi}$, we have $\mu_0 = 0$, $\mu_1 = 1/2$ and $\mu_2 = \sqrt{\frac{1}{4} - \frac{1}{2\pi}}$. It is clear that $\mathbf{f} \sim \mathcal{N}(0, \Sigma)$ with $\Sigma := \mu_1^2 \mathbf{W} \mathbf{W}^\top + \mu_2^2 \mathbf{I}$. Also the data generative model (2.1) can be written as:

$$(E.1) \quad y_i = \langle \mathbf{f}_i, \boldsymbol{\theta}_0 \rangle + w_i, \quad \text{with} \quad w_i \sim \mathcal{N}(0, \sigma^2),$$

for proper choices of σ^2 and $\boldsymbol{\theta}_0$. Indeed in both models (2.1) and (E.1), $(y_i, \mathbf{f}_i) \in \mathbb{R}^{N+1}$ is a centered Gaussian vector. By matching their covariances we obtain

$$(E.2) \quad \begin{aligned} \Sigma &= \mu_1^2 \mathbf{W} \mathbf{W}^\top + \mu_2^2 \mathbf{I}, \\ \boldsymbol{\theta}_0 &= \mu_1 \Sigma^{-1} \mathbf{W} \boldsymbol{\beta}, \\ \sigma^2 &= \tau^2 + \|\boldsymbol{\beta}\|_{\ell_2}^2 - \mu_1^2 \boldsymbol{\beta}^\top \mathbf{W}^\top \Sigma^{-1} \mathbf{W} \boldsymbol{\beta}. \end{aligned}$$

We next rewrite the objective of optimization (6.23) using this change of variable and also plug in for y_i from (E.1) to obtain

$$(E.3) \quad \mathcal{L}_{\text{nl}}(\boldsymbol{\theta}) = \frac{1}{2n} \sum_{i=1}^n (|\langle \mathbf{f}_i, \boldsymbol{\theta}_0 - \boldsymbol{\theta} \rangle + w_i| + \varepsilon \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2})^2 + \frac{\zeta}{2} \boldsymbol{\theta}^\top \boldsymbol{\Omega} \boldsymbol{\theta}.$$

We will use a powerful extension of a classical Gaussian process inequality due to Gordon [33] known as *Convex Gaussian Minimax Theorem (CGMT)* [88] to derive a precise asymptotic characterization of $\overset{\circ}{\text{AR}}_{\text{nl}}(\hat{\theta}_{\text{nl}}^*)$. A similar proof technique has been used in [46] to understand the effect of adversarial training on linear regression models. Indeed, for the particular case of $\mu_1 = 0, \mu_2 = 1$ (so $\Sigma = \mathbf{I}$) and $\mathbf{J} = \mathbf{I}$, the loss function (E.4) reduces to that studied in [46].

The CGMT analysis will output a deterministic scalar optimization which depends on ζ . We need to calculate the solution of this optimization at $\zeta \rightarrow 0$. However, as we discuss in our derivation, the objective of this optimization is strongly convex (in minimizing variables) and concave (in maximizing variables). Therefore, by continuity of its solution in the coefficients of the objective, we directly calculate the solution by setting $\zeta = 0$ in the loss $\mathcal{L}_{\text{nl}}(\boldsymbol{\theta})$, bringing us to the following restatement of the loss (with a slight abuse of notation):

$$(E.4) \quad \mathcal{L}_{\text{nl}}(\boldsymbol{\theta}) = \frac{1}{2n} \sum_{i=1}^n (|\langle \mathbf{f}_i, \boldsymbol{\theta}_0 - \boldsymbol{\theta} \rangle + w_i| + \varepsilon \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2})^2.$$

Consider a change of variable of the form $\mathbf{f}_i = \Sigma^{1/2} \mathbf{g}_i$ with $\mathbf{g}_i \sim \mathcal{N}(0, \mathbf{I})$ and $\mathbf{z} = \Sigma^{1/2}(\boldsymbol{\theta} - \boldsymbol{\theta}_0)$. Also define

$$\ell(v; \boldsymbol{\theta}) := \frac{1}{2} (|v| + \varepsilon \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2})^2.$$

Then, optimization problem (E.4) can be equivalently written in the form

$$(E.5) \quad \min_{\mathbf{z} \in \mathbb{R}^N, \mathbf{v} \in \mathbb{R}^n} \frac{1}{n} \sum_{i=1}^n \ell(v_i; \boldsymbol{\theta}_0 + \boldsymbol{\Sigma}^{-1/2} \mathbf{z}) \quad \text{subject to} \quad \mathbf{v} = \mathbf{w} - \mathbf{G}\mathbf{z}.$$

By writing the dual of this optimization problem (with dual variable $\frac{\mathbf{u}}{\sqrt{d}}$) we get

$$(E.6) \quad \begin{aligned} \min_{\mathbf{z} \in \mathbb{R}^N, \mathbf{v} \in \mathbb{R}^n} \max_{\mathbf{u} \in \mathbb{R}^n} \frac{1}{\sqrt{d}} \left\{ \mathbf{u}^\top \mathbf{G}\mathbf{z} - \mathbf{u}^\top \mathbf{w} + \mathbf{u}^\top \mathbf{v} \right\} + \frac{1}{n} \sum_{i=1}^n \ell(v_i; \boldsymbol{\theta}_0 + \boldsymbol{\Sigma}^{-1/2} \mathbf{z}) \\ = \min_{\mathbf{z} \in \mathbb{R}^N, \mathbf{v} \in \mathbb{R}^n} \max_{\mathbf{u} \in \mathbb{R}^n} \frac{1}{\sqrt{d}} \left\{ \mathbf{u}^\top \mathbf{G}\mathbf{z} - \mathbf{u}^\top \mathbf{w} + \mathbf{u}^\top \mathbf{v} \right\} + \bar{\ell}(\mathbf{v}; \mathbf{z}), \end{aligned}$$

where

$$\begin{aligned} \bar{\ell}(\mathbf{v}; \mathbf{z}) &:= \frac{1}{n} \sum_{i=1}^n \ell(v_i; \boldsymbol{\theta}_0 + \boldsymbol{\Sigma}^{-1/2} \mathbf{z}) \\ &= \frac{1}{2n} \|\mathbf{v}\|_{\ell_2}^2 + \frac{\varepsilon}{n} \|\mathbf{v}\|_{\ell_1} \|\mathbf{J}(\boldsymbol{\theta}_0 + \boldsymbol{\Sigma}^{-1/2} \mathbf{z})\|_{\ell_2} + \frac{\varepsilon^2}{2} \|\mathbf{J}(\boldsymbol{\theta}_0 + \boldsymbol{\Sigma}^{-1/2} \mathbf{z})\|_{\ell_2}^2. \end{aligned}$$

The minimax optimization (E.6) is in a form that we can apply the CGMT framework. Formally, the CGMT framework concerns problems of the form

$$(E.7) \quad \min_{\mathbf{z} \in \mathcal{S}_z} \max_{\mathbf{u} \in \mathcal{S}_u} \mathbf{u}^\top \mathbf{G}\mathbf{z} + \psi(\mathbf{z}, \mathbf{u}),$$

with $\mathbf{G}$ a matrix with i.i.d standard normal entries and shows that this problem is asymptotically equivalent to the following problem:

$$(E.8) \quad \min_{\mathbf{z} \in \mathcal{S}_z} \max_{\mathbf{u} \in \mathcal{S}_u} \|\mathbf{z}\|_{\ell_2} \mathbf{g}^\top \mathbf{u} + \|\mathbf{u}\|_{\ell_2} \mathbf{h}^\top \mathbf{z} + \psi(\mathbf{z}, \mathbf{u}),$$

where $\mathbf{g}$ and $\mathbf{h}$ are independent Gaussian vectors with i.i.d. $\mathcal{N}(0, 1)$ entries and $\psi(\mathbf{z}, \mathbf{u})$ is convex in $\mathbf{z}$ and concave in $\mathbf{u}$. Here, the sets $\mathcal{S}_z$ and $\mathcal{S}_u$ are compact sets. We refer to [88, Theorem 3] for precise statements regarding the equivalence of (E.7) and (E.8).

Following [88] we shall refer to problems of the form (E.7) as the *Primal Problem (PO)* and refer to problems of the form (E.8) as the *Auxiliary Problem (AO)*.

As described above the CGMT framework requires the minimization/maximization to be over compact sets. This technical issue can be avoided by a common trick in this literature where one introduces “artificial” boundedness constraint which do not effect the optimal solution. Specifically, following [88] we can add constraints of the form $\mathcal{S}_z = \{\mathbf{z} : \|\mathbf{z}\|_{\ell_2} \leq K_\alpha\}$ and $\mathcal{S}_u = \{\mathbf{u} : \|\mathbf{u}\|_{\ell_2} \leq K_\beta\}$ for sufficiently large constants K_α and K_β without changing the optimal solution of (E.6) in a precise asymptotic sense. We leave out a detailed argument here and refer to [46, Appendix B] for similar arguments. This allows us to replace (E.6) with

$$(E.9) \quad \min_{\mathbf{z} \in \mathcal{S}_z, \mathbf{v} \in \mathbb{R}^n} \max_{\mathbf{u} \in \mathcal{S}_u} \frac{1}{\sqrt{d}} \left\{ \mathbf{u}^\top \mathbf{G}\mathbf{z} - \mathbf{u}^\top \mathbf{w} + \mathbf{u}^\top \mathbf{v} \right\} + \bar{\ell}(\mathbf{v}; \mathbf{z}).$$

Observe that the above loss function has a bilinear term $\mathbf{u}^\top \mathbf{G}\mathbf{z}$, with $G_{ij} \sim \mathcal{N}(0, 1)$ independently, plus a function of the form

$$\psi(\mathbf{z}, \mathbf{v}, \mathbf{u}) := \frac{1}{\sqrt{d}} \left\{ -\mathbf{u}^\top \mathbf{w} + \mathbf{u}^\top \mathbf{v} \right\} + \bar{\ell}(\mathbf{v}; \mathbf{z}),$$

which is jointly convex in $(\mathbf{z}, \mathbf{v})$ and concave in $\mathbf{u}$.

Therefore the corresponding AO problem takes the following form

$$(E.10) \quad \min_{\mathbf{z} \in \mathcal{S}_z, \mathbf{v} \in \mathbb{R}^n} \max_{\mathbf{u} \in \mathcal{S}_u} \frac{1}{\sqrt{d}} \left\{ \|\mathbf{z}\|_{\ell_2} \mathbf{g}^\top \mathbf{u} + \|\mathbf{u}\|_{\ell_2} \mathbf{h}^\top \mathbf{z} - \mathbf{u}^\top \mathbf{w} + \mathbf{u}^\top \mathbf{v} \right\} + \bar{\ell}(\mathbf{v}; \mathbf{z}).$$

This concludes the derivation of the AO problem.

E.1. Scalarization of the AO problem We next simplify the AO problem by considering this problem in the asymptotic regime. We start by maximizing over $\mathbf{u}$. Write $\mathbf{u} = \beta \tilde{\mathbf{u}}$ with $\tilde{\mathbf{u}} \in \mathbb{S}^{n-1}$ and $0 \leq \beta \leq K_\beta$. Using this decomposition we have

$$\begin{aligned}
& \max_{\mathbf{u} \in \mathcal{S}_{\mathbf{u}}} \|\mathbf{z}\|_{\ell_2} \mathbf{g}^T \mathbf{u} + \|\mathbf{u}\|_{\ell_2} \mathbf{h}^T \mathbf{z} - \mathbf{u}^T \mathbf{w} + \mathbf{u}^T \mathbf{v} \\
&= \max_{0 \leq \beta \leq K_\beta} \max_{\tilde{\mathbf{u}} \in \mathbb{S}^{n-1}} \beta \|\mathbf{z}\|_{\ell_2} \mathbf{g}^T \tilde{\mathbf{u}} + \beta \mathbf{h}^T \mathbf{z} - \beta \tilde{\mathbf{u}}^T \mathbf{w} + \beta \tilde{\mathbf{u}}^T \mathbf{v} \\
&= \max_{0 \leq \beta \leq K_\beta} \max_{\tilde{\mathbf{u}} \in \mathbb{S}^{n-1}} \beta \tilde{\mathbf{u}}^T (\|\mathbf{z}\|_{\ell_2} \mathbf{g} - \mathbf{w} + \mathbf{v}) + \beta \mathbf{h}^T \mathbf{z} \\
&= \max_{0 \leq \beta \leq K_\beta} \beta \|\|\mathbf{z}\|_{\ell_2} \mathbf{g} - \mathbf{w} + \mathbf{v}\|_{\ell_2} + \beta \mathbf{h}^T \mathbf{z}.
\end{aligned}$$

After substituting the above into AO problem (E.10), it reads

$$\min_{\mathbf{z} \in \mathcal{S}_{\mathbf{z}}, \mathbf{v}} \max_{0 \leq \beta \leq K_\beta} \frac{\beta}{\sqrt{d}} \|\|\mathbf{z}\|_{\ell_2} \mathbf{g} - \mathbf{w} + \mathbf{v}\|_{\ell_2} + \frac{\beta}{\sqrt{d}} \mathbf{h}^T \mathbf{z} + \bar{\ell}(\mathbf{v}; \mathbf{z}).$$

We next aim to simplify the minimization over $\mathbf{v}$ and $\mathbf{z}$, but a hurdle is that they are coupled through the term $\bar{\ell}(\mathbf{v}; \mathbf{z})$. To address this technical issue, we consider the conjugate of $\bar{\ell}(\mathbf{v}; \mathbf{z})$ in with respect to $\mathbf{z}$. That is,

$$\bar{\ell}(\mathbf{v}; \mathbf{z}) = \sup_{\mathbf{q}} \mathbf{q}^T \mathbf{z} - \tilde{\ell}(\mathbf{v}; \mathbf{q}).$$

The AO problem then can be written as

$$\text{(E.11)} \quad \min_{\mathbf{z} \in \mathcal{S}_{\mathbf{z}}, \mathbf{v}} \max_{0 \leq \beta \leq K_\beta, \mathbf{q}} \frac{\beta}{\sqrt{d}} \|\|\mathbf{z}\|_{\ell_2} \mathbf{g} - \mathbf{w} + \mathbf{v}\|_{\ell_2} + \frac{\beta}{\sqrt{d}} \mathbf{h}^T \mathbf{z} + \mathbf{q}^T \mathbf{z} - \tilde{\ell}(\mathbf{v}; \mathbf{q}).$$

In the above optimization the order of minimization and maximization can be flipped using the Sion's theorem and the fact that the original PO problem is convex/concave in the min/max parameters. The argument only uses the convexity of the loss $\ell(\mathbf{v}; \mathbf{q})$ and we refer to [86, Appendix A.2.4] for a detailed argument. This brings us to

$$\max_{0 \leq \beta \leq K_\beta, \mathbf{q}} \min_{\mathbf{z} \in \mathcal{S}_{\mathbf{z}}, \mathbf{v}} \frac{\beta}{\sqrt{d}} \|\|\mathbf{z}\|_{\ell_2} \mathbf{g} - \mathbf{w} + \mathbf{v}\|_{\ell_2} + \frac{\beta}{\sqrt{d}} \mathbf{h}^T \mathbf{z} + \mathbf{q}^T \mathbf{z} - \tilde{\ell}(\mathbf{v}; \mathbf{q}).$$

We optimize over the direction and norm of $\mathbf{z}$ ($\|\mathbf{z}\|_{\ell_2} = \alpha$) to get

$$\text{(E.12)} \quad \max_{0 \leq \beta \leq K_\beta, \mathbf{q}} \min_{0 \leq \alpha \leq K_\alpha, \mathbf{v}} \frac{\beta}{\sqrt{d}} \|\alpha \mathbf{g} - \mathbf{w} + \mathbf{v}\|_{\ell_2} - \alpha \left\| \frac{\beta}{\sqrt{d}} \mathbf{h} + \mathbf{q} \right\|_{\ell_2} - \tilde{\ell}(\mathbf{v}; \mathbf{q}).$$

Note that $\tilde{\ell}(\mathbf{v}; \mathbf{q})$ is convex in $\mathbf{q}$ and so the AO objective (E.12) is clearly jointly concave in $\mathbf{q}$ and β . Also since $\bar{\ell}$ is jointly convex in $(\mathbf{v}, \mathbf{z})$, then $-\bar{\ell}(\mathbf{v}; \mathbf{z})$ is jointly concave in $(\mathbf{v}, \mathbf{z})$. Also $\mathbf{q}^T \mathbf{z}$ is jointly concave in $(\mathbf{v}, \mathbf{z})$. Therefore, $\mathbf{q}^T \mathbf{z} - \bar{\ell}(\mathbf{v}; \mathbf{z})$ is jointly concave in $(\mathbf{v}, \mathbf{z})$ and based on the partial maximization rule we can conclude that $\tilde{\ell}(\mathbf{v}; \mathbf{q})$ should be concave in $\mathbf{v}$. The other terms are also trivially jointly convex in $\alpha, \mathbf{v}$ so that overall the objective is jointly convex in $\alpha, \mathbf{v}$. Therefore, by virtue of Sion's min-max Theorem [76]) we can change the order of the mins and maxs as we please and rewrite the AO problem as

$$\min_{0 \leq \alpha \leq K_\alpha, \mathbf{v}} \max_{0 \leq \beta \leq K_\beta, \mathbf{q}} \frac{\beta}{\sqrt{d}} \|\alpha \mathbf{g} - \mathbf{w} + \mathbf{v}\|_{\ell_2} - \alpha \left\| \frac{\beta}{\sqrt{d}} \mathbf{h} + \mathbf{q} \right\|_{\ell_2} - \tilde{\ell}(\mathbf{v}; \mathbf{q}).$$

To continue we shall calculate the conjugate function $\tilde{\ell}$. This is the subject of the next lemma.

Lemma E.1 *The conjugate of*

$$\bar{\ell}(\mathbf{v}; \mathbf{z}) := \frac{1}{2n} \sum_{i=1}^n \left(|v_i| + \varepsilon \|\mathbf{J}(\boldsymbol{\theta}_0 + \boldsymbol{\Sigma}^{-1/2} \mathbf{z})\|_{\ell_2} \right)^2,$$

with respect to the variable $\mathbf{z}$ is given by

$$\tilde{\ell}(\mathbf{v}; \mathbf{q}) := \sup_{\mathbf{z}} \mathbf{q}^T \mathbf{z} - \bar{\ell}(\mathbf{v}; \mathbf{z}) = -\langle \boldsymbol{\Sigma}^{1/2} \boldsymbol{\theta}_0, \mathbf{q} \rangle + \frac{1}{2} \left(\frac{1}{\varepsilon} \|\boldsymbol{\Sigma}^{1/2} \mathbf{J}^{-1} \mathbf{q}\|_{\ell_2} - \frac{1}{n} \|\mathbf{v}\|_{\ell_1} \right)_+^2 - \frac{1}{2n} \|\mathbf{v}\|_{\ell_2}^2.$$

We refer to Section E.5 for the proof of this lemma. Plugging in for $\tilde{\ell}$ from Lemma E.1 in the AO problem we arrive at

$$\begin{aligned} \min_{0 \leq \alpha < K_\alpha, \mathbf{v}} \max_{0 \leq \beta \leq K_\beta, \mathbf{q}} & \frac{\beta}{\sqrt{d}} \|\alpha \mathbf{g} - \mathbf{w} + \mathbf{v}\|_{\ell_2} - \alpha \left\| \frac{\beta}{\sqrt{d}} \mathbf{h} + \mathbf{q} \right\|_{\ell_2} \\ & + \langle \boldsymbol{\Sigma}^{1/2} \boldsymbol{\theta}_0, \mathbf{q} \rangle - \frac{1}{2} \left(\frac{1}{\varepsilon} \|\mathbf{J}^{-1} \boldsymbol{\Sigma}^{1/2} \mathbf{q}\|_{\ell_2} - \frac{\|\mathbf{v}\|_{\ell_1}}{n} \right)_+^2 + \frac{1}{2n} \|\mathbf{v}\|_{\ell_2}^2. \end{aligned} \quad (\text{E.13})$$

Optimization over $\mathbf{q}$: To simplify the AO problem further, we next focus on maximization over $\mathbf{q}$. Consider the change of variable $\tilde{\mathbf{q}} := \mathbf{J}^{-1} \boldsymbol{\Sigma}^{1/2} \mathbf{q}$ and keep only the terms in the AO objective which involve $\mathbf{q}$.

$$\begin{aligned} & \max_{\mathbf{q}} -\alpha \left\| \frac{\beta}{\sqrt{d}} \mathbf{h} + \mathbf{q} \right\|_{\ell_2} + \langle \boldsymbol{\Sigma}^{1/2} \boldsymbol{\theta}_0, \mathbf{q} \rangle - \frac{1}{2} \left(\frac{1}{\varepsilon} \|\mathbf{J}^{-1} \boldsymbol{\Sigma}^{1/2} \mathbf{q}\|_{\ell_2} - \frac{\|\mathbf{v}\|_{\ell_1}}{n} \right)_+^2 \\ & = \max_{\tilde{\mathbf{q}}} -\alpha \left\| \frac{\beta}{\sqrt{d}} \mathbf{h} + \boldsymbol{\Sigma}^{-1/2} \mathbf{J} \tilde{\mathbf{q}} \right\|_{\ell_2} + \langle \boldsymbol{\theta}_0, \mathbf{J} \tilde{\mathbf{q}} \rangle - \frac{1}{2} \left(\frac{1}{\varepsilon} \|\tilde{\mathbf{q}}\|_{\ell_2} - \frac{\|\mathbf{v}\|_{\ell_1}}{n} \right)_+^2 \\ & = \max_{\tilde{\mathbf{q}}, 0 \leq \tau_q} -\frac{\alpha}{2\tau_q} \left\| \frac{\beta}{\sqrt{d}} \mathbf{h} + \boldsymbol{\Sigma}^{-1/2} \mathbf{J} \tilde{\mathbf{q}} \right\|_{\ell_2}^2 - \frac{\alpha \tau_q}{2} + \langle \boldsymbol{\theta}_0, \mathbf{J} \tilde{\mathbf{q}} \rangle - \frac{1}{2} \left(\frac{1}{\varepsilon} \|\tilde{\mathbf{q}}\|_{\ell_2} - \frac{\|\mathbf{v}\|_{\ell_1}}{n} \right)_+^2 \\ & = \max_{\tilde{\mathbf{q}}, 0 \leq \tau_q} -\frac{\alpha}{2\tau_q} \left\| \frac{\beta}{\sqrt{d}} \mathbf{h} + \boldsymbol{\Sigma}^{-1/2} \mathbf{J} \tilde{\mathbf{q}} - \frac{\tau_q}{\alpha} \boldsymbol{\Sigma}^{1/2} \boldsymbol{\theta}_0 \right\|_{\ell_2}^2 + \frac{\tau_q}{2\alpha} \|\boldsymbol{\Sigma}^{1/2} \boldsymbol{\theta}_0\|_{\ell_2}^2 - \left\langle \frac{\beta}{\sqrt{d}} \mathbf{h}, \boldsymbol{\Sigma}^{1/2} \boldsymbol{\theta}_0 \right\rangle \\ & \quad - \frac{\alpha \tau_q}{2} - \frac{1}{2} \left(\frac{1}{\varepsilon} \|\tilde{\mathbf{q}}\|_{\ell_2} - \frac{\|\mathbf{v}\|_{\ell_1}}{n} \right)_+^2 \end{aligned}$$

We next maximize over $\tilde{\mathbf{q}}$ by introducing a new dummy variable γ for $\|\tilde{\mathbf{q}}\|_{\ell_2}$. This brings us to the following problem

$$\begin{aligned} (\text{E.14}) \quad & \min_{\tilde{\mathbf{q}}, 0 \leq \gamma} \frac{\alpha}{2\tau_q} \left\| \frac{\beta}{\sqrt{d}} \mathbf{h} + \boldsymbol{\Sigma}^{-1/2} \mathbf{J} \tilde{\mathbf{q}} - \frac{\tau_q}{\alpha} \boldsymbol{\Sigma}^{1/2} \boldsymbol{\theta}_0 \right\|_{\ell_2}^2 + \frac{1}{2} \left(\frac{1}{\varepsilon} \gamma - \frac{\|\mathbf{v}\|_{\ell_1}}{n} \right)_+^2 \\ & \text{subject to } \|\tilde{\mathbf{q}}\|_{\ell_2} = \gamma. \end{aligned}$$

We continue by the following lemma and refer to Section E.5 for its proof.

Lemma E.2 *Let $H \in \mathbb{R}^{d \times d}$ be invertible and $\mathbf{r} \in \mathbb{R}^d$, $c_0, c_1 \in \mathbb{R}$. Consider the following optimization problem:*

$$\begin{aligned} (\text{E.15}) \quad & \min_{\tilde{\mathbf{q}}, 0 \leq \gamma} \frac{c_0}{2} \|\mathbf{H} \tilde{\mathbf{q}} - \mathbf{r}\|_{\ell_2}^2 + \frac{1}{2} \left(\frac{1}{\varepsilon} \gamma - c_1 \right)_+^2 \\ & \text{s.t. } \|\tilde{\mathbf{q}}\|_{\ell_2} = \gamma \end{aligned}$$

Define

$$(E.16) \quad Q(\mathbf{H}, \mathbf{r}, \gamma) = \sup_{\lambda \geq 0} \frac{\lambda}{2} (\mathbf{r}^\top (\mathbf{H}\mathbf{H}^\top + \lambda \mathbf{I})^{-1} \mathbf{r} - \gamma^2).$$

Then, the optimal objective value of (E.15) is given by

$$\min_{\gamma \geq 0} c_0 Q(\mathbf{H}, \mathbf{r}, \gamma) + \frac{1}{2} \left(\frac{1}{\varepsilon} \gamma - c_1 \right)_+^2.$$

Using Lemma E.2, the optimal value of (E.14) is given by

$$\min_{\gamma \geq 0} \frac{\alpha}{\tau_q} Q(\Sigma^{-1/2} \mathbf{J}, \frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h}, \gamma) + \frac{1}{2} \left(\frac{\gamma}{\varepsilon} - \frac{\|\mathbf{v}\|_{\ell_1}}{n} \right)_+^2.$$

Therefore, by substituting in (E.12) the AO optimization can be simplified as

$$(E.17) \quad \min_{0 \leq \alpha < K_\alpha, \mathbf{v}} \max_{0 \leq \beta \leq K_\beta, 0 \leq \gamma, \tau_q} \frac{\beta}{\sqrt{d}} \|\alpha \mathbf{g} - \mathbf{w} + \mathbf{v}\|_{\ell_2} - \frac{\alpha}{\tau_q} Q(\Sigma^{-1/2} \mathbf{J}, \frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h}, \gamma) \\ + \frac{\tau_q}{2\alpha} \|\Sigma^{1/2} \boldsymbol{\theta}_0\|_{\ell_2}^2 - \langle \frac{\beta}{\sqrt{d}} \mathbf{h}, \Sigma^{1/2} \boldsymbol{\theta}_0 \rangle - \frac{\alpha \tau_q}{2} - \frac{1}{2} \left(\frac{\gamma}{\varepsilon} - \frac{\|\mathbf{v}\|_{\ell_1}}{n} \right)_+^2 + \frac{1}{2n} \|\mathbf{v}\|_{\ell_2}^2.$$

Before proceeding further with our simplification of the AO problem, let us state the following lemma which is used to discuss the convexity-concavity of the objective and justification of changing the order of maximization and minimization. We refer to Section E.5 for its proof.

Lemma E.3 *The function*

$$f(\gamma, \beta, \tau_q) := \frac{1}{\tau_q} Q(\Sigma^{-1/2} \mathbf{J}, \frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h}, \gamma),$$

is jointly convex in the variables $(\gamma, \frac{\beta}{\sqrt{d}}, \tau_q)$.

As a result this lemma, the objective (E.17) is jointly concave in (γ, β, τ_q) . Also recall that since $\tilde{\ell}$ was concave the objective (E.12) was jointly convex in $(\alpha, \mathbf{v})$. Since maximization (with respect to direction of $\tilde{\mathbf{q}}$) preserves convexity (pointwise maximum of convex functions is convex), therefore the objective (E.17) is jointly convex in $(\alpha, \mathbf{v})$.

Therefore, by another use of Sion's min-max theorem, we can change the order of min and max in (E.17) and write it equivalently as

$$(E.18) \quad \max_{0 \leq \beta \leq K_\beta, 0 \leq \gamma, \tau_q} \min_{0 \leq \alpha < K_\alpha, \mathbf{v}} \frac{\beta}{\sqrt{d}} \|\alpha \mathbf{g} - \mathbf{w} + \mathbf{v}\|_{\ell_2} - \frac{\alpha}{\tau_q} Q(\Sigma^{-1/2} \mathbf{J}, \frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h}, \gamma) \\ + \frac{\tau_q}{2\alpha} \|\Sigma^{1/2} \boldsymbol{\theta}_0\|_{\ell_2}^2 - \frac{1}{\sqrt{d}} \langle \beta \mathbf{h}, \Sigma^{1/2} \boldsymbol{\theta}_0 \rangle - \frac{\alpha \tau_q}{2} - \frac{1}{2} \left(\frac{\gamma}{\varepsilon} - \frac{\|\mathbf{v}\|_{\ell_1}}{n} \right)_+^2 + \frac{1}{2n} \|\mathbf{v}\|_{\ell_2}^2.$$

We next focus on minimization over $\mathbf{v}$. Keeping only the terms in (E.17) that depend on $\mathbf{v}$ we have

$$(E.19) \quad \min_{\mathbf{v}} \frac{\beta}{\sqrt{d}} \|\alpha \mathbf{g} - \mathbf{w} + \mathbf{v}\|_{\ell_2} + \frac{1}{2n} \|\mathbf{v}\|_{\ell_2}^2 - \frac{1}{2} \left(\frac{\gamma}{\varepsilon} - \frac{\|\mathbf{v}\|_{\ell_1}}{n} \right)_+^2 \\ = \min_{\tau_g \geq 0, \mathbf{v}} \frac{\beta}{2n\tau_g} \|\alpha \mathbf{g} - \mathbf{w} + \mathbf{v}\|_{\ell_2}^2 + \frac{\beta \tau_g n}{2d} + \frac{1}{2n} \|\mathbf{v}\|_{\ell_2}^2 - \frac{1}{2n^2} \left(\frac{n\gamma}{\varepsilon} - \|\mathbf{v}\|_{\ell_1} \right)_+^2.$$

Recall the definition of the Moreau envelope function of a function f at a point $\mathbf{x}$ with parameter μ ,

$$e_f(\mathbf{x}; \rho) \equiv \min_{\mathbf{v}} \frac{1}{2\rho} \|\mathbf{x} - \mathbf{v}\|_{\ell_2}^2 + f(\mathbf{v}).$$

and define

$$(E.20) \quad f(\mathbf{v}; \gamma) := \frac{1}{2} \|\mathbf{v}\|_{\ell_2}^2 - \frac{1}{2n} \left(\frac{n}{\varepsilon} \gamma - \|\mathbf{v}\|_{\ell_1} \right)_+^2.$$

Note that $f(\mathbf{v}; \gamma)$ is convex in $\mathbf{v}$ (since $-\tilde{\ell}(\mathbf{v}; \mathbf{q})$ was convex in $\mathbf{v}$). Thus, (E.19) can be rewritten in the more compact form

$$(E.21) \quad \min_{\tau_g \geq 0} \frac{1}{n} e_f \left(\mathbf{w} - \alpha \mathbf{g}; \frac{\tau_g}{\beta} \right) + \frac{\beta \tau_g}{2} \frac{n}{d}.$$

We next invoke the result of [46, Lemma 6.3] which gives a characterization of the Moreau envelope function $e_f(\mathbf{x}; \mu)$.

Lemma E.4 ([46, Lemma 6.3]) *Consider the function f given by (E.20). Then,*

$$e_f(\mathbf{x}; \rho) = \frac{1}{2(\rho+1)} \|\mathbf{x}\|_{\ell_2}^2 + \min_{\nu \geq 0} G_n(\mathbf{x}; \rho, \gamma, \nu),$$

where

$$(E.22) \quad G_n(\mathbf{x}; \rho, \gamma, \nu) = \frac{1}{2\rho(\rho+1)} \|\mathbf{x} - \text{ST}(\mathbf{x}; \nu)\|_{\ell_2}^2 - \frac{1}{2n} \left(\frac{n}{\varepsilon} \gamma - \frac{1}{1+\rho} \|\text{ST}(\mathbf{x}; \nu)\|_{\ell_1} \right)_+^2,$$

and $\text{ST}(\mathbf{x}; \nu)$ is the soft-thresholding function defined as

$$[\text{ST}(\mathbf{x}; \nu)]_i = \begin{cases} x_i - \lambda, & \text{if } x_i \geq \lambda, \\ 0 & \text{if } |x_i| \leq \lambda, \\ x_i + \lambda & \text{if } x_i \leq -\lambda, \end{cases}$$

for each coordinate i . Furthermore, $e_f(\mathbf{x}; \tau)$ is strictly convex in $\mathbf{x}$.

Using this characterization in (E.19) we get

$$(E.23) \quad \begin{aligned} & \min_{\mathbf{v}} \frac{\beta}{\sqrt{d}} \|\alpha \mathbf{g} - \mathbf{w} + \mathbf{v}\|_{\ell_2} + \frac{1}{2n} \|\mathbf{v}\|_{\ell_2}^2 - \frac{1}{2} \left(\frac{\gamma}{\varepsilon} - \frac{\|\mathbf{v}\|_{\ell_1}}{n} \right)_+^2 \\ &= \min_{\tau_g \geq 0} \frac{1}{n} e_f \left(\mathbf{w} - \alpha \mathbf{g}; \frac{\tau_g}{\beta} \right) + \frac{\beta \tau_g}{2} \frac{n}{d} \\ &= \min_{\tau_g \geq 0} \frac{\beta \tau_g}{2} \frac{n}{d} + \frac{1}{n} \frac{\beta}{2(\tau_g + \beta)} \|\mathbf{w} - \alpha \mathbf{g}\|_{\ell_2}^2 + \frac{1}{n} \min_{\nu \geq 0} G_n(\mathbf{w} - \alpha \mathbf{g}; \frac{\tau_g}{\beta}, \gamma, \nu). \end{aligned}$$

Next by plugging (E.23) in (E.18) we arrive at the following AO formulation:

$$(E.24) \quad \begin{aligned} & \max_{0 \leq \beta \leq K_\beta, 0 \leq \gamma, \tau_q} \min_{0 \leq \alpha < K_\alpha, 0 \leq \tau_g, \nu} - \frac{\alpha}{\tau_q} Q(\Sigma^{-1/2} \mathbf{J}, \frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h}, \gamma) + \frac{\tau_q}{2\alpha} \|\Sigma^{1/2} \boldsymbol{\theta}_0\|_{\ell_2}^2 - \frac{\alpha \tau_q}{2} \\ & - \frac{1}{\sqrt{d}} \langle \beta \mathbf{h}, \Sigma^{1/2} \boldsymbol{\theta}_0 \rangle + \frac{\beta \tau_g}{2} \frac{n}{d} + \frac{\beta}{2(\tau_g + \beta)} \frac{1}{n} \|\mathbf{w} - \alpha \mathbf{g}\|_{\ell_2}^2 \end{aligned}$$

$$(E.25) \quad + \frac{1}{n} G_n(\mathbf{w} - \alpha \mathbf{g}; \frac{\tau_g}{\beta}, \gamma, \nu).$$

Recall that the problem (E.19) was jointly convex in $(\mathbf{v}, \alpha, \tau_g)$ and (E.18) jointly concave in (β, γ, τ_q) . Since partial minimization preserves convexity we therefore conclude that the objective (E.24) is jointly convex in (α, τ_g) and jointly concave in (β, γ, τ_q) (after the minimization over $\nu \geq 0$ has been carried out).

E.2. Convergence analysis of the AO problem

E.2.1. Pointwise convergence We next derive the pointwise limit of the AO objective in the asymptotic regime that $N/d \rightarrow \psi_1$ and $n/d \rightarrow \psi_2$, as $n \rightarrow \infty$.

Recalling the definition of $\boldsymbol{\theta}_0$ from (E.2) we have

$$\|\Sigma^{1/2}\boldsymbol{\theta}_0\|_{\ell_2}^2 = \mu_1^2 \|\Sigma^{-1/2}\mathbf{W}\boldsymbol{\beta}\|_{\ell_2}^2 = \frac{\mu_1^2 \|\boldsymbol{\beta}\|_{\ell_2}^2}{d} \text{trace}(\mathbf{W}^\top \Sigma \mathbf{W}),$$

where we used the fact that the distribution of $\mathbf{W}$ is rotationally invariant. By our assumption $\|\boldsymbol{\beta}\|_{\ell_2} \rightarrow 1$. Let $0 \leq s_1, \dots, s_N$ denote the eigenvalues of $\mathbf{W}\mathbf{W}^\top$. By invoking the definition of Σ from (E.2) we have

$$\begin{aligned} \|\Sigma^{1/2}\boldsymbol{\theta}_0\|_{\ell_2}^2 &= \frac{\psi_1 \mu_1^2 \|\boldsymbol{\beta}\|_{\ell_2}^2}{N} \text{trace}(\mathbf{W}^\top \Sigma^{-1} \mathbf{W}) \\ &= \frac{\psi_1 \mu_1^2 \|\boldsymbol{\beta}\|_{\ell_2}^2}{N} \sum_{i=1}^d \frac{s_i}{\mu_1^2 s_i + \mu_2^2} \\ (E.26) \quad &= \frac{\psi_1 \mu_1^2 \|\boldsymbol{\beta}\|_{\ell_2}^2}{N} \sum_{i=1}^d \frac{1}{\mu_1^2} \left(1 - \frac{\mu_2^2 / \mu_1^2}{s_i + \mu_2^2 / \mu_1^2} \right) \rightarrow \psi_1 \left(1 + \frac{\mu_2^2}{\mu_1^2} S\left(-\frac{\mu_2^2}{\mu_1^2}; \psi_1\right) \right), \end{aligned}$$

in probability where $S(z) = \int \frac{\rho(s)}{z-s} ds$ is the Stieltjes transform of the spectral density ρ of the matrix $\mathbf{W}\mathbf{W}^\top$. The formula for $S(z)$ is given in Proposition F.2 and since it is a function of ψ_1 , we make this dependence clear in the notation and write $S(z; \psi_1)$ henceforth.

Since for our activation $\mu_1 = \frac{1}{2}$ and $\mu_2 = \sqrt{\frac{1}{4} - \frac{1}{2\pi}}$, this simplifies to

$$(E.27) \quad \|\Sigma^{1/2}\boldsymbol{\theta}_0\|_{\ell_2}^2 \rightarrow \psi_1 \left(1 + \left(1 - \frac{2}{\pi} \right) S\left(\frac{2}{\pi} - 1; \psi_1\right) \right),$$

in probability. This together with (E.2) implies that

$$(E.28) \quad \sigma^2 = \tau^2 + \|\boldsymbol{\beta}\|_{\ell_2}^2 - \|\Sigma^{1/2}\boldsymbol{\theta}_0\|_{\ell_2}^2 \rightarrow \tau^2 + 1 - \psi_1 \left(1 + \left(1 - \frac{2}{\pi} \right) S\left(\frac{2}{\pi} - 1; \psi_1\right) \right).$$

We next note that since $\|\Sigma^{1/2}\boldsymbol{\theta}_0\|_{\ell_2} = O(1)$, for $\mathbf{h} \sim \mathcal{N}(0, \mathbf{I})$ we have

$$(E.29) \quad \frac{1}{\sqrt{n}} \langle \mathbf{h}, \Sigma^{1/2}\boldsymbol{\theta}_0 \rangle \rightarrow 0,$$

in probability, as $n \rightarrow \infty$. In addition, since $\mathbf{g} \sim \mathcal{N}(0, \mathbf{I}_n)$ and $\mathbf{w} \sim \mathcal{N}(0, \sigma^2 \mathbf{I}_n)$, we have

$$(E.30) \quad \frac{1}{n} \|\alpha \mathbf{g} - \mathbf{w}\|_{\ell_2}^2 \rightarrow \alpha^2 + \sigma^2,$$

in probability.

We next proceed by calculating the limit of the Q function. By definition,

$$\begin{aligned} &Q(\Sigma^{-1/2}\mathbf{J}, \frac{\tau_q}{\alpha} \Sigma^{1/2}\boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}}\mathbf{h}, \gamma) \\ &= \sup_{\lambda \geq 0} \frac{\lambda}{2} \left[\left(\frac{\tau_q}{\alpha} \Sigma^{1/2}\boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}}\mathbf{h} \right)^\top (\Sigma^{-1/2}\mathbf{J}^2 \Sigma^{-1/2} + \lambda \mathbf{I})^{-1} \left(\frac{\tau_q}{\alpha} \Sigma^{1/2}\boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}}\mathbf{h} \right) - \gamma^2 \right] \\ (E.31) \quad &= \sup_{\lambda \geq 0} \frac{\lambda}{2} \left[\left(\frac{\tau_q}{\alpha} \Sigma \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \Sigma^{1/2} \mathbf{h} \right)^\top (\mathbf{J}^2 + \lambda \Sigma)^{-1} \left(\frac{\tau_q}{\alpha} \Sigma \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \Sigma^{1/2} \mathbf{h} \right) - \gamma^2 \right]. \end{aligned}$$

We compute the limit of the right hand side for any fixed value of $\lambda \geq 0$. First note that since $\|(\Sigma^{-1/2} \mathbf{J}^2 \Sigma^{-1/2} + \lambda \mathbf{I})^{-1}\| = O_p(1)$ and by invoking (E.29), the cross terms vanish in the limit and we have

$$\begin{aligned} & \lim_{n \rightarrow \infty} \left(\frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h} \right)^\top (\Sigma^{-1/2} \mathbf{J}^2 \Sigma^{-1/2} + \lambda \mathbf{I})^{-1} \left(\frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h} \right) \\ (E.32) \quad & = \lim_{n \rightarrow \infty} \frac{\tau_q^2}{\alpha^2} \boldsymbol{\theta}_0^\top \Sigma^{1/2} (\Sigma^{-1/2} \mathbf{J}^2 \Sigma^{-1/2} + \lambda \mathbf{I})^{-1} \Sigma^{1/2} \boldsymbol{\theta}_0 + \lim_{n \rightarrow \infty} \frac{\beta^2}{d} \mathbf{h}^\top (\Sigma^{-1/2} \mathbf{J}^2 \Sigma^{-1/2} + \lambda \mathbf{I})^{-1} \mathbf{h}. \end{aligned}$$

We treat each term separately. Plugging for $\boldsymbol{\theta}_0$ from (E.2) we have

$$\begin{aligned} & \lim_{n \rightarrow \infty} \frac{\tau_q^2}{\alpha^2} \boldsymbol{\theta}_0^\top \Sigma^{1/2} (\Sigma^{-1/2} \mathbf{J}^2 \Sigma^{-1/2} + \lambda \mathbf{I})^{-1} \Sigma^{1/2} \boldsymbol{\theta}_0 = \lim_{n \rightarrow \infty} \left(\frac{\mu_1 \tau_q}{\alpha} \right)^2 \boldsymbol{\beta}^\top \mathbf{W}^\top (\mathbf{J}^2 + \lambda \Sigma)^{-1} \mathbf{W} \boldsymbol{\beta} \\ (E.33) \quad & = \lim_{n \rightarrow \infty} \left(\frac{\mu_1 \tau_q}{\alpha} \right)^2 \frac{1}{d} \text{trace}(\mathbf{W}^\top (\mathbf{J}^2 + \lambda \Sigma)^{-1} \mathbf{W}), \end{aligned}$$

where in the last step we used the fact that the distribution of $\mathbf{W}$ is rotationally invariant and $\|\boldsymbol{\beta}\|_{\ell_2} \rightarrow 1$.

Similarly since $\mathbf{h} \sim \mathcal{N}(0, \mathbf{I}_N)$ we have

$$\begin{aligned} & \lim_{n \rightarrow \infty} \frac{\beta^2}{d} \mathbf{h}^\top (\Sigma^{-1/2} \mathbf{J}^2 \Sigma^{-1/2} + \lambda \mathbf{I})^{-1} \mathbf{h} = \lim_{n \rightarrow \infty} \frac{\beta^2}{d} \langle \Sigma^{1/2} (\mathbf{J}^2 + \lambda \Sigma)^{-1} \Sigma^{1/2}, \mathbf{h} \mathbf{h}^\top \rangle \\ (E.34) \quad & = \lim_{n \rightarrow \infty} \frac{\beta^2}{d} \text{trace}(\Sigma^{1/2} (\mathbf{J}^2 + \lambda \Sigma)^{-1} \Sigma^{1/2}). \end{aligned}$$

Combining (E.33) and (E.34) into (E.32) we get

$$\begin{aligned} & \lim_{n \rightarrow \infty} \left(\frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h} \right)^\top (\Sigma^{-1/2} \mathbf{J}^2 \Sigma^{-1/2} + \lambda \mathbf{I})^{-1} \left(\frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h} \right) \\ (E.35) \quad & = \lim_{n \rightarrow \infty} \text{trace} \left\{ (\mathbf{J}^2 + \lambda \Sigma)^{-1} \left(\left(\frac{\mu_1 \tau_q}{\alpha} \right)^2 \frac{1}{d} \mathbf{W} \mathbf{W}^\top + \frac{\beta^2}{d} \Sigma \right) \right\}, \end{aligned}$$

where we used that $\text{trace}(\mathbf{A}\mathbf{B}) = \text{trace}(\mathbf{B}\mathbf{A})$ for any two matrices $\mathbf{A}$ and $\mathbf{B}$.

To calculate the limit on the right-hand side of (E.35), we use Proposition F.3 on the spectrum of inner product kernel random matrices.

Since $\|\mathbf{W}\| = O_p(1)$ we also have $\left\| \left(\frac{\mu_1 \tau_q}{\alpha} \right)^2 \frac{1}{d} \mathbf{W} \mathbf{W}^\top + \frac{\beta^2}{d} \Sigma \right\| = O_p(1)$ and as an immediate corollary of Proposition F.3, in (E.35) we can replace $\mathbf{J}^2$ with $\mathbf{K}$ since they have the same spectrum. This brings us to

$$\begin{aligned} & \lim_{n \rightarrow \infty} \left(\frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h} \right)^\top (\Sigma^{-1/2} \mathbf{J}^2 \Sigma^{-1/2} + \lambda \mathbf{I})^{-1} \left(\frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h} \right) \\ (E.36) \quad & = \lim_{n \rightarrow \infty} \text{trace} \left\{ (\mathbf{K} + \lambda \Sigma)^{-1} \left(\left(\frac{\mu_1 \tau_q}{\alpha} \right)^2 \cdot \frac{1}{d} \mathbf{W} \mathbf{W}^\top + \frac{\beta^2}{d} \Sigma \right) \right\} \\ & = \lim_{n \rightarrow \infty} \frac{1}{d} \text{trace} \left\{ \left(\left(\frac{1}{4} + \lambda \mu_1^2 \right) \mathbf{W} \mathbf{W}^\top + \left(\frac{1}{4} + \lambda \mu_2^2 \right) \mathbf{I} \right)^{-1} \left(\left(\frac{\mu_1^2 \tau_q^2}{\alpha^2} + \beta^2 \mu_1^2 \right) \mathbf{W} \mathbf{W}^\top + \beta^2 \mu_2^2 \mathbf{I} \right) \right\}. \end{aligned}$$

Note that the latter only depends on the spectral density of $\mathbf{W} \mathbf{W}^\top$ and can be written in terms of its Stieltjes transform.

By law of large numbers and with simple algebraic manipulations it is easy to see that for any constants b_0, b_1, c_0, c_1 we have

$$(E.37) \quad \frac{1}{N} \sum_{i=1}^N \frac{b_0 s_i + b_1}{c_0 s_i + c_1} \rightarrow \frac{b_0}{c_0} + \frac{b_0 \frac{c_1}{c_0} - b_1}{c_0} S(-c_1/c_0; \psi_1),$$

with $S(t; \psi_1)$ representing the Stieltjes transform of the spectral density of $\mathbf{W}\mathbf{W}^\top$.

Using this with (E.36) we obtain

$$(E.38) \quad \lim_{n \rightarrow \infty} \left(\frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h} \right)^\top (\Sigma^{-1/2} \mathbf{J}^2 \Sigma^{-1/2} + \lambda \mathbf{I})^{-1} \left(\frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h} \right) \\ = \frac{4\psi_1}{1+4\lambda\mu_1^2} \left(\frac{\mu_1^2 \tau_q^2}{\alpha^2} + \beta^2 \mu_1^2 \right) + \frac{4\psi_1}{1+4\lambda\mu_1^2} \left\{ \left(\frac{\mu_1^2 \tau_q^2}{\alpha^2} + \beta^2 \mu_1^2 \right) \left(\frac{1+4\lambda\mu_2^2}{1+4\lambda\mu_1^2} \right) - \beta^2 \mu_2^2 \right\} S\left(-\frac{1+4\lambda\mu_2^2}{1+4\lambda\mu_1^2}; \psi_1\right).$$

By combining (E.38) and (E.31) we get

$$(E.39) \quad \lim_{n \rightarrow \infty} Q(\Sigma^{-1/2} \mathbf{J}, \frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{n}} \mathbf{h}, \gamma) \\ = \sup_{\lambda \geq 0} \frac{\lambda}{2} \left[\frac{4\psi_1}{1+4\lambda\mu_1^2} \left(\frac{\mu_1^2 \tau_q^2}{\alpha^2} + \beta^2 \mu_1^2 \right) + \frac{4\psi_1}{1+4\lambda\mu_1^2} \left\{ \left(\frac{\mu_1^2 \tau_q^2}{\alpha^2} + \beta^2 \mu_1^2 \right) \left(\frac{1+4\lambda\mu_2^2}{1+4\lambda\mu_1^2} \right) - \beta^2 \mu_2^2 \right\} S\left(-\frac{1+4\lambda\mu_2^2}{1+4\lambda\mu_1^2}; \psi_1\right) - \gamma^2 \right].$$

Plugging for $\mu_1 = \frac{1}{2}$ and $\mu_2 = \sqrt{\frac{1}{4} - \frac{1}{2\pi}}$ we have

$$(E.40) \quad \lim_{n \rightarrow \infty} Q(\Sigma^{-1/2} \mathbf{J}, \frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{n}} \mathbf{h}, \gamma) = F\left(\frac{\tau_q}{\alpha}, \beta, \psi_1, \gamma\right),$$

with the definition

$$(E.41) \quad F(a, b, \psi_1, \gamma) := \sup_{\lambda \geq 0} \frac{\lambda \psi_1}{2(1+\lambda)} \left\{ a^2 + b^2 + \left(a^2 \left(1 - \frac{2}{\pi} \frac{\lambda}{1+\lambda} \right) + \frac{2b^2}{\pi(1+\lambda)} \right) S\left(\frac{2}{\pi} \frac{\lambda}{1+\lambda} - 1; \psi_1\right) \right\} - \frac{\lambda}{2} \gamma^2.$$

By the change of variable $\tilde{\lambda} = \frac{\lambda}{1+\lambda}$, the function F can be written as

$$(E.42) \quad F(a, b, \psi_1, \gamma) := \sup_{0 \leq \tilde{\lambda} < 1} \frac{\tilde{\lambda} \psi_1}{2} \left\{ a^2 + b^2 + \left(a^2 \left(1 - \frac{2}{\pi} \tilde{\lambda} \right) + \frac{2(1-\tilde{\lambda})b^2}{\pi} \right) S\left(\frac{2}{\pi} \tilde{\lambda} - 1; \psi_1\right) \right\} - \frac{\tilde{\lambda}}{2(1-\tilde{\lambda})} \gamma^2.$$

We next proceed to characterize the limit of $\frac{1}{n} G_n(\mathbf{w} - \alpha \mathbf{g}; \frac{\tau_g}{\beta}, \gamma, \nu)$. To this end, we recall the result of [46, Lemma 6.4].

Lemma E.5 *Let $\mathbf{u} \in \mathbb{R}^n$ be a Gaussian random vector distributed as $\mathcal{N}(\mathbf{0}, \omega^2 \mathbf{I}_n)$. Then,*

(E.43)

$$\lim_{n \rightarrow \infty} \frac{1}{2n\rho(\rho+1)} \|\mathbf{u} - \text{ST}(\mathbf{u}; \nu)\|_{\ell_2}^2 = \frac{\omega^2}{2\rho(\rho+1)} \left(\left(1 - \sqrt{\frac{2}{\pi}} \frac{\nu}{\omega} e^{-\frac{\nu^2}{2\omega^2}} \right) \right)^2,$$

(E.44)

$$\lim_{n \rightarrow \infty} \frac{1}{2n^2} \left(\frac{n}{\varepsilon} \gamma - \frac{1}{1+\rho} \|\text{ST}(\mathbf{u}; \nu)\|_{\ell_1} \right)_+^2 = \frac{\omega^2}{2(\rho+1)^2} \left(\frac{\gamma(\rho+1)}{\varepsilon\omega} + \frac{\nu}{\omega} \cdot \text{erfc}\left(\frac{1}{\sqrt{2}} \frac{\nu}{\omega}\right) - \sqrt{\frac{2}{\pi}} e^{-\frac{\nu^2}{2\omega^2}} \right)_+^2.$$

Therefore, by (E.22) we have

$$\begin{aligned} \lim_{n \rightarrow \infty} \frac{1}{n} G_n(\mathbf{u}; \rho, \gamma, \nu) &= \frac{\omega^2}{2\rho(\rho+1)} \left(\left(1 - \sqrt{\frac{2}{\pi}} \frac{\nu}{\omega} e^{-\frac{\nu^2}{2\omega^2}} \right) + \left(\frac{\nu^2}{\omega^2} - 1 \right) \operatorname{erfc} \left(\frac{1}{\sqrt{2}} \frac{\nu}{\omega} \right) \right) \\ &\quad - \frac{\omega^2}{2(\rho+1)^2} \left(\frac{\gamma(\rho+1)}{\varepsilon\omega} + \frac{\nu}{\omega} \cdot \operatorname{erfc} \left(\frac{1}{\sqrt{2}} \frac{\nu}{\omega} \right) - \sqrt{\frac{2}{\pi}} e^{-\frac{\nu^2}{2\omega^2}} \right)^2. \end{aligned}$$

Furthermore,

$$\begin{aligned} \min_{\nu \geq 0} \lim_{n \rightarrow \infty} \frac{1}{n} G_n(\mathbf{u}; \rho, \gamma, \nu) \\ = G(\omega; \rho, \gamma) := \begin{cases} 0 & \text{if } \gamma(\rho+1) \leq \sqrt{\frac{2}{\pi}} \varepsilon\omega \\ \frac{\omega^2}{2\rho(\rho+1)} \left(\operatorname{erf} \left(\frac{\nu^* \left(\frac{\gamma(\rho+1)}{\varepsilon\omega}, \rho \right)}{\sqrt{2}} \right) - \frac{\gamma(\rho+1)}{\varepsilon\omega} \nu^* \left(\frac{\gamma(\rho+1)}{\varepsilon\omega}, \rho \right) \right) & \text{if } \gamma(\rho+1) > \sqrt{\frac{2}{\pi}} \varepsilon\omega \end{cases} \end{aligned}$$

where $\nu^*(a, \rho)$ is the unique solution to

$$a - \frac{1}{\rho} \nu - \nu \cdot \operatorname{erf} \left(\frac{\nu}{\sqrt{2}} \right) - \sqrt{\frac{2}{\pi}} e^{-\frac{\nu^2}{2}} = 0.$$

Combining (E.27), (E.29), (E.30), (E.40), and Lemma E.5, we obtain the following scalarized AO problem:

$$\begin{aligned} \max_{0 \leq \beta \leq K_\beta, 0 \leq \gamma, \tau_q} \min_{0 \leq \alpha < K_\alpha, 0 \leq \tau_g} & -\frac{\alpha}{\tau_q} F\left(\frac{\tau_q}{\alpha}, \beta, \psi_1, \gamma\right) + \frac{\tau_q}{2\alpha} (\tau^2 + 1 - \sigma^2) - \frac{\alpha \tau_q}{2} \\ & + \frac{\beta \tau_g}{2} \psi_2 + \frac{\beta}{2(\tau_g + \beta)} (\sigma^2 + \alpha^2) + G\left(\sqrt{\sigma^2 + \alpha^2}; \frac{\tau_g}{\beta}, \gamma\right), \end{aligned} \quad (\text{E.45})$$

where $\sigma^2 = \tau^2 + 1 - \psi_1 \left(1 + \left(1 - \frac{2}{\pi} \right) S\left(\frac{2}{\pi} - 1; \psi_1\right) \right)$.

We conclude this part by a lemma on the convexity-concavity of the above scalarized AO problem and the uniqueness of the solution to the AO problem.

Lemma E.6 (Strict convexity and uniqueness of the solution) *The objective function (E.45) is strictly jointly convex in (α, τ_g) and jointly concave in (β, γ, τ_q) . Also the solution $(\alpha_*, \frac{\tau_{g*}}{\beta_*})$ to this problem is unique.*

We defer the proof of Lemma E.6 to Section E.5. This concludes the proof of Theorem 4.2(a).

E.2.2. Uniform convergence In Section E.2.1 we showed that the objective function in (E.24) converges point-wise to the objective function in (E.45). However, for our goal we need to show that the minimax solutions of the converging sequence of the objectives in (E.24) converges to the minimax solution of the AO objective in (E.45), denoted by $\mathcal{R}(\alpha, \tau_g, \beta, \gamma, \tau_q)$. Convexity/concavity of $\mathcal{R}$ plays a crucial role here since it is being used to conclude local uniform convergence from the point-wise convergence.

This can be shown by following similar arguments as in [86, Lemma A.5] that is essentially based on a result known as “convexity lemma” in the literature (see e.g. [54, Lemma 7.75]) by which point-wise convergence of convex functions, of a finite number of variables, implies uniform convergence in compact subsets. Since the argument here is general, we leave out a detailed discussion and refer to [86, Lemma A.5].

E.3. Proof of Theorem 4.2(b) In Proposition 6.8 we gave a characterization of $\overset{\circ}{\text{AR}}_{\text{nl}}$. We first provide an alternative characterization in terms of the equivalent model of (E.2).

Recall the key quantity a from Proposition 6.8, given by

$$(E.46) \quad a^2 = \tau^2 + \left\| \frac{1}{2} \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta} \right\|_{\ell_2}^2 + \left(\frac{1}{4} - \frac{1}{2\pi} \right) \|\boldsymbol{\theta}\|_{\ell_2}^2.$$

We claim that $a^2 = \sigma^2 + \|\Sigma^{1/2}(\boldsymbol{\theta} - \boldsymbol{\theta}_0)\|_{\ell_2}^2$. To see this, we expand this expression as follows:

$$\begin{aligned} \sigma^2 + \|\Sigma^{1/2}(\boldsymbol{\theta} - \boldsymbol{\theta}_0)\|_{\ell_2}^2 &= \sigma^2 + \langle \boldsymbol{\theta}, \Sigma \boldsymbol{\theta} \rangle + \langle \boldsymbol{\theta}_0, \Sigma \boldsymbol{\theta}_0 \rangle - 2\langle \boldsymbol{\theta}_0, \Sigma \boldsymbol{\theta} \rangle \\ &= \sigma^2 + \mu_1^2 \|\mathbf{W}^\top \boldsymbol{\theta}\|_{\ell_2}^2 + \mu_2^2 \|\boldsymbol{\theta}\|_{\ell_2}^2 + \mu_1^2 \boldsymbol{\beta}^\top \mathbf{W}^\top \Sigma^{-1} \mathbf{W} \boldsymbol{\beta} - 2\mu_1 \langle \mathbf{W}, \boldsymbol{\beta}, \boldsymbol{\theta} \rangle \\ &= \sigma^2 + \|\mu_1 \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta}\|_{\ell_2}^2 - \|\boldsymbol{\beta}\|_{\ell_2}^2 + \mu_2^2 \|\boldsymbol{\theta}\|_{\ell_2}^2 + \mu_1^2 \boldsymbol{\beta}^\top \mathbf{W}^\top \Sigma^{-1} \mathbf{W} \boldsymbol{\beta} \\ &= \tau^2 + \|\mu_1 \mathbf{W}^\top \boldsymbol{\theta} - \boldsymbol{\beta}\|_{\ell_2}^2 + \mu_2^2 \|\boldsymbol{\theta}\|_{\ell_2}^2, \end{aligned}$$

where we used the definition of Σ , $\boldsymbol{\theta}_0$ and σ^2 as per (E.2). The claim follows by recalling that for the shifted Relu activation, $\mu_1 = \frac{1}{2}$ and $\mu_2 = \sqrt{\frac{1}{4} - \frac{1}{2\pi}}$.

By the above characterization of quantity a we obtain

$$(E.47) \quad a^2 = \sigma^2 + \|\Sigma^{1/2}(\boldsymbol{\theta} - \boldsymbol{\theta}_0)\|_{\ell_2}^2.$$

We next note that by definition of the variables in the AO problem, we have $\mathbf{z} = \Sigma^{1/2}(\boldsymbol{\theta} - \boldsymbol{\theta}_0)$ and $\alpha = \|\mathbf{z}\|_{\ell_2}$. Therefore,

$$\lim_{n \rightarrow \infty} \|\Sigma^{1/2}(\boldsymbol{\theta} - \boldsymbol{\theta}_0)\|_{\ell_2} = \alpha_*.$$

Invoking the limit of σ^2 given by (E.28), we get

$$(E.48) \quad \lim_{n \rightarrow \infty} a^2 = \tau^2 + 1 - \psi_1 \left(1 + \left(1 - \frac{2}{\pi} \right) S \left(\frac{2}{\pi} - 1 \right) \right) + \alpha_*^2.$$

We next characterize $\lim_{n \rightarrow \infty} \|\mathbf{J}\boldsymbol{\theta}\|_{\ell_2}$. We will use the same AO problem to calculate this quantity.

Recall that $\widehat{\mathbf{z}} = \Sigma^{1/2}(\widehat{\boldsymbol{\theta}}_{\text{nl}}^* - \boldsymbol{\theta}_0)$ satisfies the following relation with $\mathbf{q}_*$ the optimizer in (E.13):

$$\mathbf{q}_* = \arg \max_{\mathbf{q}} \mathbf{q}^\top \widehat{\mathbf{z}} - \widetilde{\ell}(\mathbf{v}; \mathbf{q}),$$

where $\widetilde{\ell}(\mathbf{v}; \mathbf{q})$ is the convex conjugate of $\bar{\ell}(\mathbf{v}; \mathbf{z})$. Since conjugate of a conjugate function is the function itself we then have

$$\begin{aligned} \widehat{\mathbf{z}} &= \arg \max_{\mathbf{z}} \mathbf{q}_*^\top \mathbf{z} - \bar{\ell}(\mathbf{v}; \mathbf{z}) \\ (E.49) \quad &= \arg \max_{\mathbf{z}} \mathbf{q}_*^\top \mathbf{z} - \frac{1}{2n} \|\mathbf{v}\|_{\ell_2}^2 - \frac{\varepsilon}{n} \|\mathbf{v}\|_{\ell_1} \|\mathbf{J}(\boldsymbol{\theta}_0 + \Sigma^{-1/2} \mathbf{z})\|_{\ell_2} - \frac{\varepsilon^2}{2} \|\mathbf{J}(\boldsymbol{\theta}_0 + \Sigma^{-1/2} \mathbf{z})\|_{\ell_2}^2. \end{aligned}$$

We consider two cases:

Case 1: $\|\mathbf{J}(\boldsymbol{\theta}_0 + \Sigma^{-1/2} \widehat{\mathbf{z}})\|_{\ell_2} \neq 0$. Setting derivative with respect to $\widehat{\mathbf{z}}$ to zero we obtain

$$\mathbf{q}_* - \frac{\varepsilon}{n} \|\mathbf{v}\|_{\ell_1} \frac{\Sigma^{-1/2} \mathbf{J}^2(\boldsymbol{\theta}_0 + \Sigma^{-1/2} \widehat{\mathbf{z}})}{\|\mathbf{J}(\boldsymbol{\theta}_0 + \Sigma^{-1/2} \widehat{\mathbf{z}})\|_{\ell_2}} - \varepsilon^2 \Sigma^{-1/2} \mathbf{J}^2(\boldsymbol{\theta}_0 + \Sigma^{-1/2} \widehat{\mathbf{z}}) = 0.$$

By rearranging the terms we write it as

$$\mathbf{J}(\boldsymbol{\theta}_0 + \boldsymbol{\Sigma}^{-1/2} \hat{\mathbf{z}}) = \left(\frac{\varepsilon}{n} \frac{\|\mathbf{v}\|_{\ell_1}}{\|\mathbf{J}(\boldsymbol{\theta}_0 + \boldsymbol{\Sigma}^{-1/2} \hat{\mathbf{z}})\|_{\ell_2}} + \varepsilon^2 \right)^{-1} \mathbf{J}^{-1} \boldsymbol{\Sigma}^{1/2} \mathbf{q}_*.$$

By taking the ℓ_2 norm of both sides and then solving for $\|\mathbf{J}(\boldsymbol{\theta}_0 + \boldsymbol{\Sigma}^{-1/2} \hat{\mathbf{z}})\|_{\ell_2}$, we get

$$(E.50) \quad \|\mathbf{J}(\boldsymbol{\theta}_0 + \boldsymbol{\Sigma}^{-1/2} \hat{\mathbf{z}})\|_{\ell_2} = \frac{1}{\varepsilon^2} \|\mathbf{J}^{-1} \boldsymbol{\Sigma}^{1/2} \mathbf{q}_*\|_{\ell_2} - \frac{1}{n\varepsilon} \|\mathbf{v}\|_{\ell_1}.$$

Case 2: $\|\mathbf{J}(\boldsymbol{\theta}_0 + \boldsymbol{\Sigma}^{-1/2} \hat{\mathbf{z}})\|_{\ell_2} = 0$. In this case, $\hat{\mathbf{z}} = -\boldsymbol{\Sigma}^{1/2} \boldsymbol{\theta}_0$ and by comparing the objective function of (E.49) at the optimal solution under case 1 and case 2, it is easy to verify that case 2 happens only when the right-hand side in (E.50) becomes negative. Therefore, the two cases can be combined together in the following form:

$$(E.51) \quad \langle \hat{\boldsymbol{\theta}}_{\text{nl}}^*, \mathbf{J}^2 \hat{\boldsymbol{\theta}}_{\text{nl}}^* \rangle = \|\mathbf{J}(\boldsymbol{\theta}_0 + \boldsymbol{\Sigma}^{-1/2} \hat{\mathbf{z}})\|_{\ell_2}^2 = \left(\frac{1}{\varepsilon^2} \|\mathbf{J}^{-1} \boldsymbol{\Sigma}^{1/2} \mathbf{q}_*\|_{\ell_2} - \frac{1}{n\varepsilon} \|\mathbf{v}\|_{\ell_1} \right)_+^2.$$

So in order to get the asymptotic value of the left hand side we can work with the right-hand side with $\mathbf{v}$ and $\gamma = \|\mathbf{J}^{-1} \boldsymbol{\Sigma}^{1/2} \mathbf{q}_*\|_{\ell_2}$ the optimal solutions of the AO problem.

In Lemma E.4 (which is a restatement of [46, Lemma 6.3]), the Moreau envelop function $e_f(\mathbf{x}, \rho)$ was characterized. Following the proof of [46, Lemma 6.3], we can verify that the optimal $\mathbf{v}$ is given by

$$\mathbf{v} = \frac{1}{1 + \frac{\tau_g}{\beta}} \text{ST}(\mathbf{w} - \alpha \mathbf{g}; \nu).$$

Therefore, by invoking the relation (E.44) we have

$$\lim_{n \rightarrow \infty} \frac{1}{n^2} \left(\frac{n}{\varepsilon} \gamma - \|\mathbf{v}\|_{\ell_1} \right)_+^2 = \frac{\omega^2}{(\rho + 1)^2} \left(\frac{\gamma(\rho + 1)}{\varepsilon \omega} + \nu_* \cdot \text{erfc} \left(\frac{1}{\sqrt{2}} \nu_* \right) - \sqrt{\frac{2}{\pi}} e^{-\nu_*^2} \right)_+^2,$$

with $\omega = \sqrt{\alpha^2 + \sigma^2}$, $\rho = \frac{\tau_g}{\beta}$, $\nu_* = \nu_*(\frac{\gamma(\rho+1)}{\varepsilon \omega}, \rho)$ and $\nu_*(a, \mu)$ the unique solution to the following equation:

$$a - \frac{1}{\rho} \nu - \nu \cdot \text{erf} \left(\frac{\nu}{\sqrt{2}} \right) - \sqrt{\frac{2}{\pi}} e^{-\frac{\nu^2}{2}} = 0.$$

Plugging the above relation into (E.51) we obtain

$$\begin{aligned} \lim_{n \rightarrow \infty} \langle \hat{\boldsymbol{\theta}}_{\text{nl}}^*, \mathbf{J}^2 \hat{\boldsymbol{\theta}}_{\text{nl}}^* \rangle &= \lim_{n \rightarrow \infty} \left(\frac{1}{\varepsilon^2} \gamma - \frac{1}{n\varepsilon} \|\mathbf{v}\|_{\ell_1} \right)_+^2 \\ &= \frac{1}{\varepsilon^2 n^2} \lim_{n \rightarrow \infty} \left(\frac{n}{\varepsilon} \gamma - \|\mathbf{v}\|_{\ell_1} \right)_+^2 \\ &= \frac{\omega^2}{\varepsilon^2 (\rho + 1)^2} \left(\frac{\gamma(\rho + 1)}{\varepsilon \omega} + \nu_* \cdot \text{erfc} \left(\frac{1}{\sqrt{2}} \nu_* \right) - \sqrt{\frac{2}{\pi}} e^{-\nu_*^2} \right)_+^2 \\ &= \frac{\omega^2}{\varepsilon^2 (\rho + 1)^2} \left(\frac{\gamma(\rho + 1)}{\varepsilon \omega} + \frac{\rho + 1}{\rho} \nu_* - \frac{\gamma(\rho + 1)}{\varepsilon \omega} \right)_+^2 \\ &= \frac{\omega^2}{\varepsilon^2 (\rho + 1)^2} \left(\frac{1 + \rho}{\rho} \nu_* \right)_+^2 \end{aligned}$$

$$\begin{aligned}
&= \frac{\omega^2 \nu_*^2}{\varepsilon^2 \rho^2} \\
&= \frac{(\alpha_*^2 + \sigma^2) \nu_*^2}{\varepsilon^2 \rho^2} \\
(E.52) \quad &= \frac{\beta^2 \nu_*^2 (\alpha_*^2 + \sigma^2)}{\varepsilon^2 \tau_g^2}.
\end{aligned}$$

Now by recalling the characterization (6.20) along with (E.52) and (E.47) we get the desired result of (4.6).

E.4. Proof of Proposition 5.1 Recall the objective $\mathcal{R}(\alpha, \tau_g, \beta, \gamma, \tau_q)$. We start by considering the change of variable $\tilde{\gamma} = \gamma/\varepsilon$. Note that the term $-\lambda/(1-\lambda)\gamma^2 = -\lambda/(1-\lambda)\varepsilon^2\tilde{\gamma}^2$ will be dropped as it is zero. We next argue that $\nu^* = 0$. The reason is that if the indicator in the objective is inactive then the corresponding term is void, which is equivalent to $\nu^* = 0$. If the indicator is active, since the problem is maximization over γ , we deduce that $\frac{\gamma(\tau_g + \beta)}{\varepsilon\beta\sqrt{\alpha^2 + \sigma^2}} = \sqrt{\frac{2}{\pi}}$, which by the equation defining ν^* implies that $\nu^* = 0$. Next, by straightforward calculation, and using definition of Stieltjes transform S , we have that the expression inside $\sup_{0 \leq \lambda < 1}$ is increasing in λ and so we have that the optimal $\lambda \rightarrow 1$. Using these values, the objective reduces to

$$\begin{aligned}
\mathcal{R}(\alpha, \tau_g, \beta, \tau_q) &= \frac{\tau_q}{2\alpha}(\tau^2 + 1 - \sigma^2) - \frac{\alpha\tau_q}{2} + \frac{\beta\tau_g}{2}\psi_2 + \frac{\beta}{2(\tau_g + \beta)}(\sigma^2 + \alpha^2) \\
&\quad - \frac{\psi_1}{2} \left\{ \frac{\tau_q}{\alpha} + \frac{\alpha}{\tau_q}\beta^2 + \frac{\tau_q}{\alpha} \left(1 - \frac{2}{\pi}\right) S\left(\frac{2}{\pi} - 1; \psi_1\right) \right\}.
\end{aligned}$$

Using the definition

$$\sigma^2 = \tau^2 + 1 - \psi_1 \left(1 + \left(1 - \frac{2}{\pi}\right) S\left(\frac{2}{\pi} - 1; \psi_1\right)\right),$$

we can further simplify the objective as

$$\mathcal{R}(\alpha, \tau_g, \beta, \tau_q) = -\frac{\psi_1}{2} \frac{\beta^2 \alpha}{\tau_q} - \frac{\alpha\tau_q}{2} + \frac{\beta\tau_g}{2}\psi_2 + \frac{\beta}{2(\tau_g + \beta)}(\sigma^2 + \alpha^2).$$

Optimization over τ_q can be done easily resulting in $\tau_q = \beta\sqrt{\psi_1}$, which gives

$$\mathcal{R}(\alpha, \tau_g, \beta) = -\alpha\beta\sqrt{\psi_1} + \frac{\beta\tau_g}{2}\psi_2 + \frac{\beta}{2(\tau_g + \beta)}(\sigma^2 + \alpha^2).$$

Writing the stationary condition for α, τ_g, β we arrive at the following system of equations:

$$(E.53) \quad \begin{cases} \frac{\alpha}{\tau_g + \beta} = \sqrt{\psi_1}, \\ \psi_2 = \frac{\sigma^2 + \alpha^2}{(\tau_g + \beta)^2}, \\ -\alpha\sqrt{\psi_1} + \frac{\tau_g\psi_2}{2} + \frac{\tau_g}{2} \frac{\sigma^2 + \alpha^2}{(\tau_g + \beta)^2} = 0. \end{cases}$$

Solving the above system of equations we obtain $\alpha^2 = \sigma^2\psi_1/(\psi_2 - \psi_1)$. Recalling that ν^* , using Theorem 4.2 (b) we get the standard risk of the estimator to be

$$\text{SR}(\hat{\theta}) = \alpha_*^2 + \sigma^2 = \sigma^2 \left(\frac{\psi_2}{\psi_2 - \psi_1} \right).$$

E.5. Proofs of the Auxiliary Lemmas

E.5.1. *Proof of Lemma E.1* We start by considering the following related but different function

$$\ell_0(\mathbf{v}; \mathbf{z}) = \frac{1}{2n} \sum_{i=1}^n \left(|v_i| + \varepsilon \|\mathbf{z}\|_{\ell_2} \right)^2.$$

As shown in the proof of Lemma 6.1 in [46], the conjugate of this function is given by

$$\ell_0^*(\mathbf{v}; \mathbf{q}) = \frac{1}{2} \left(\frac{\|\mathbf{q}\|_{\ell_2}}{\varepsilon} - \frac{\|\mathbf{v}\|_{\ell_1}}{n} \right)_+^2 - \frac{1}{2n} \|\mathbf{v}\|_{\ell_2}^2.$$

Note that $\bar{\ell}(\mathbf{v}; \mathbf{z}) = \ell_0(\mathbf{v}; \mathbf{J}(\boldsymbol{\theta}_0 + \boldsymbol{\Sigma}^{-1/2} \mathbf{z}))$. We next use the result that if $f(\mathbf{x}) = g(\mathbf{A}\mathbf{x} + \mathbf{x}_0)$ then the conjugate of f can be written in terms of the conjugate of g as follows:

$$f^*(\mathbf{y}) = -\langle \mathbf{A}^{-1} \mathbf{x}_0, \mathbf{y} \rangle + g^*(\mathbf{A}^{-\top} \mathbf{y}).$$

Using this result with $\mathbf{x}_0 = \mathbf{J}\boldsymbol{\theta}_0$ and $\mathbf{A} = \mathbf{J}\boldsymbol{\Sigma}^{-1/2}$ we obtain

$$\bar{\ell}(\mathbf{v}; \mathbf{q}) = -\langle \boldsymbol{\Sigma}^{1/2} \boldsymbol{\theta}_0, \mathbf{q} \rangle + \frac{1}{2} \left(\frac{1}{\varepsilon} \|\mathbf{J}^{-1} \boldsymbol{\Sigma}^{1/2} \mathbf{q}\|_{\ell_2} - \frac{\|\mathbf{v}\|_{\ell_1}}{n} \right)_+^2 - \frac{1}{2n} \|\mathbf{v}\|_{\ell_2}^2$$

E.5.2. *Proof of Lemma E.2* Consider a slightly different optimization than (E.15) where the equality constraint is replaced by the inequality constraint $\|\tilde{\mathbf{q}}\|_{\ell_2} \leq \gamma$:

$$(E.54) \quad \min_{\tilde{\mathbf{q}}, 0 \leq \gamma} \frac{c_0}{2} \|\mathbf{H}\tilde{\mathbf{q}} - \mathbf{r}\|_{\ell_2}^2 + \frac{1}{2} \left(\frac{1}{\varepsilon} \gamma - c_1 \right)_+^2$$

s.t. $\|\tilde{\mathbf{q}}\|_{\ell_2} \leq \gamma.$

Due to this change, (E.54) is now a convex optimization. Denote by OPT_1 the optimal objective value of the original problem (E.15) and by OPT_2 the optimal objective value of the modified problem (E.54). We argue that $\text{OPT}_1 = \text{OPT}_2$. Clearly $\text{OPT}_1 \geq \text{OPT}_2$ because (E.54) has a larger feasible set. Now suppose that this inequality is strict ($\text{OPT}_1 > \text{OPT}_2$) and let $(\tilde{\mathbf{q}}_*, \gamma_*)$ be a solution to (E.54). Then we should have $\|\tilde{\mathbf{q}}_*\|_{\ell_2} < \gamma_*$. Consider the point $(\tilde{\mathbf{q}}_*, \|\tilde{\mathbf{q}}_*\|_{\ell_2})$ which is a feasible point for both optimization problems and so the objective value at this point is at least OPT_1 and therefore strictly larger than OPT_2 . But this is a contradiction because $\left(\frac{1}{\varepsilon} \gamma - c_1 \right)_+^2$ is non-decreasing in $\gamma \geq 0$.

To characterize OPT_2 , we first focus on the minimization over $\tilde{\mathbf{q}}$. The corresponding Lagrangian with Lagrange multiplier $\frac{\lambda c_0}{2}$ reads

$$\sup_{\lambda \geq 0} \min_{\tilde{\mathbf{q}}} \frac{c_0}{2} \|\mathbf{H}\tilde{\mathbf{q}} - \mathbf{r}\|_{\ell_2}^2 - \frac{\lambda c_0}{2} (\gamma^2 - \|\tilde{\mathbf{q}}\|_{\ell_2}^2).$$

Solving the inner minimization, we have $\tilde{\mathbf{q}}_* = (\mathbf{H}^\top \mathbf{H} + \lambda \mathbf{I})^{-1} \mathbf{H}^\top \mathbf{r}$ and the dual problem becomes

$$\begin{aligned} & \sup_{\lambda \geq 0} \frac{c_0}{2} \|\mathbf{H}\tilde{\mathbf{q}}_* - \mathbf{r}\|_{\ell_2}^2 - \frac{\lambda c_0}{2} (\gamma^2 - \|\tilde{\mathbf{q}}_*\|_{\ell_2}^2) \\ &= \sup_{\lambda \geq 0} \frac{c_0}{2} \tilde{\mathbf{q}}_*^\top [\mathbf{H}^\top (\mathbf{H}\tilde{\mathbf{q}}_* - \mathbf{r}) + \lambda \tilde{\mathbf{q}}_*] - \frac{c_0}{2} \mathbf{r}^\top (\mathbf{H}\tilde{\mathbf{q}}_* - \mathbf{r}) - \frac{\lambda c_0}{2} \gamma^2 \\ &= \sup_{\lambda \geq 0} -\frac{c_0}{2} \mathbf{r}^\top (\mathbf{H}\tilde{\mathbf{q}}_* - \mathbf{r}) - \frac{\lambda c_0}{2} \gamma^2 \end{aligned}$$

$$\begin{aligned}
&= \sup_{\lambda \geq 0} -\frac{c_0}{2} \mathbf{r}^\top (\mathbf{H}(\mathbf{H}^\top \mathbf{H} + \lambda \mathbf{I})^{-1} \mathbf{H}^\top - \mathbf{I}) \mathbf{r} - \frac{\lambda c_0}{2} \gamma^2 \\
&= \sup_{\lambda \geq 0} \frac{\lambda c_0}{2} \mathbf{r}^\top (c_0 \mathbf{H}^\top \mathbf{H} + \lambda \mathbf{I})^{-1} \mathbf{r} - \frac{\lambda c_0}{2} \gamma^2 \\
\text{(E.55)} \quad &= c_0 Q(\mathbf{H}, \mathbf{r}, \gamma).
\end{aligned}$$

By the Slater's condition the duality gap is zero and hence by next minimizing over $\gamma \geq 0$, we obtain that the optimal value of (E.54) is given by

$$\min_{\gamma \geq 0} c_0 Q(\mathbf{H}, \mathbf{r}, \gamma) + \frac{1}{2} \left(\frac{1}{\varepsilon} \gamma - c_1 \right)_+^2.$$

E.5.3. Proof of Lemma E.3 We first show that the function

$$g(\gamma, \beta) = Q(\Sigma^{-1/2} \mathbf{J}, \frac{1}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h}, \gamma)$$

is jointly convex in (γ, β) . By the change of variable $\tilde{\lambda} = \lambda\gamma$, $\tilde{\boldsymbol{\theta}} = \Sigma^{1/2} \boldsymbol{\theta}_0 / \alpha$, $\mathbf{H} = \Sigma^{-1/2} \mathbf{J}$, this function can be written as

$$\begin{aligned}
g(\gamma, \beta) &= \sup_{\tilde{\lambda} \geq 0} \frac{\tilde{\lambda}}{2\gamma} \left(\left(\tilde{\boldsymbol{\theta}} - \frac{\beta}{\sqrt{d}} \mathbf{h} \right)^\top (\mathbf{H} \mathbf{H}^\top + \frac{\tilde{\lambda}}{\gamma} \mathbf{I})^{-1} \left(\tilde{\boldsymbol{\theta}} - \frac{\beta}{\sqrt{d}} \mathbf{h} \right) - \gamma^2 \right) \\
&= \sup_{\tilde{\lambda} \geq 0} \frac{\tilde{\lambda}}{2} \left(\left(\tilde{\boldsymbol{\theta}} - \frac{\beta}{\sqrt{d}} \mathbf{h} \right)^\top (\gamma \mathbf{H} \mathbf{H}^\top + \tilde{\lambda} \mathbf{I})^{-1} \left(\tilde{\boldsymbol{\theta}} - \frac{\beta}{\sqrt{d}} \mathbf{h} \right) - \gamma \right).
\end{aligned}$$

We show that for any fixed $\tilde{\lambda} \geq 0$ the inner function above is jointly convex in (γ, β) and since the pointwise maximum of convex functions is also convex, we conclude that $g(\gamma, \beta)$ is jointly convex in (γ, β) .

The Hessian of the inner function reads

$$\frac{1}{2} \nabla_{\frac{\beta}{\sqrt{d}}, \gamma}^2 \left[\left(\tilde{\boldsymbol{\theta}} - \frac{\beta}{\sqrt{d}} \mathbf{h} \right)^\top (\gamma \mathbf{H} \mathbf{H}^\top + \tilde{\lambda} \mathbf{I})^{-1} \left(\tilde{\boldsymbol{\theta}} - \frac{\beta}{\sqrt{d}} \mathbf{h} \right) - \gamma \right] = \begin{bmatrix} A & C \\ C & B \end{bmatrix},$$

where

$$\begin{aligned}
A &:= \mathbf{h}^\top (\gamma \mathbf{H} \mathbf{H}^\top + \tilde{\lambda} \mathbf{I})^{-1} \mathbf{h} \\
B &:= \left\| \left(\gamma \mathbf{H} \mathbf{H}^\top + \tilde{\lambda} \mathbf{I} \right)^{-1/2} \mathbf{H} \mathbf{H}^\top (\gamma \mathbf{H} \mathbf{H}^\top + \tilde{\lambda} \mathbf{I})^{-1} \left(\tilde{\boldsymbol{\theta}} - \frac{\beta}{\sqrt{d}} \mathbf{h} \right) \right\|_{\ell_2}^2 \\
C &:= \mathbf{h}^\top (\gamma \mathbf{H} \mathbf{H}^\top + \tilde{\lambda} \mathbf{I})^{-1} \mathbf{H} \mathbf{H}^\top (\gamma \mathbf{H} \mathbf{H}^\top + \tilde{\lambda} \mathbf{I})^{-1} \left(\tilde{\boldsymbol{\theta}} - \frac{\beta}{\sqrt{d}} \mathbf{h} \right).
\end{aligned}$$

Here we repeatedly used the identity $\frac{\partial \mathbf{K}^{-1}}{\partial \gamma} = -\mathbf{K}^{-1} \frac{\partial \mathbf{K}}{\partial \gamma} \mathbf{K}^{-1}$, for a matrix $\mathbf{K}$.

To lighten the notation, we set $\mathbf{M} := (\gamma \mathbf{H} \mathbf{H}^\top + \tilde{\lambda} \mathbf{I})^{-1}$ and $\mathbf{v} := \tilde{\boldsymbol{\theta}} - \frac{\beta}{\sqrt{d}} \mathbf{h}$. Using these shorthands the determinant of the Hessian is equal to

$$\| \mathbf{M}^{1/2} \mathbf{h} \|_{\ell_2}^2 \| \mathbf{M}^{1/2} \mathbf{H} \mathbf{H}^\top \mathbf{M} \mathbf{v} \|_{\ell_2}^2 - (\mathbf{h}^\top \mathbf{M} \mathbf{H} \mathbf{H}^\top \mathbf{M} \mathbf{v})^2 \geq 0,$$

using the Cauchy–Schwarz inequality. This completes the proof of $g(\gamma, \beta)$ being jointly convex in (γ, β) .

Next note that its perspective function is given by

$$\begin{aligned}\tau_q g(\gamma/\tau_q, \beta/\tau_q) &= \tau_q Q(\Sigma^{-1/2} \mathbf{J}, \frac{1}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\tau_q \sqrt{d}} \mathbf{h}, \frac{\gamma}{\tau_q}) \\ &= \frac{1}{\tau_q} Q(\Sigma^{-1/2} \mathbf{J}, \frac{\tau_q}{\alpha} \Sigma^{1/2} \boldsymbol{\theta}_0 - \frac{\beta}{\sqrt{d}} \mathbf{h}, \gamma),\end{aligned}$$

and therefore is jointly convex in (γ, β, τ_q) .

E.5.4. Proof of Lemma E.6 As we discussed after (E.24), the objective function in (E.24) is jointly convex in (α, τ_g) and jointly concave in (β, γ, τ_q) . Since convexity/concavity is preserved by point-wise limits, the objective (E.45) is jointly convex in (α, τ_g) and jointly concave in (β, γ, τ_q) . To prove strict convexity in (α, τ_g) , note that in our derivation we wrote (E.19) (the part of the objective E.18 that involves $\mathbf{v}$) in terms of the Moreau envelope $\frac{1}{n} e_f(\mathbf{w} - \alpha \mathbf{g}; \frac{\tau_g}{\beta})$, cf. (E.21). As $d \rightarrow \infty$, its limit goes to the *expected Moreau envelope*. By using the result of [86, Lemma 4.4] the expected Moreau envelope of a function is strictly convex in $\mathbb{R}_{>0} \times \mathbb{R}_{>0}$ without requiring any strong or strict convexity assumption on the function itself. Therefore, the objective (E.24) (and so objective of (E.45) after taking point-wise limit) is jointly strictly convex in (α, τ_g) .

To prove the uniqueness, note that $\max_{0 \leq \beta, \gamma, \tau_q} \mathcal{R}(\alpha, \tau_g, \beta, \gamma, \tau_q)$ is strictly convex in (α, τ_g) . This follows from the fact that if $f(\mathbf{x}, \mathbf{y})$ is strictly convex in $\mathbf{x}$, then $\max_{\mathbf{y}} f(\mathbf{x}, \mathbf{y})$ is also strictly convex in $\mathbf{x}$. We next use [86, Lemma C.5] to conclude that $\min_{\tau_g > 0} \max_{0 \leq \beta, \gamma, \tau_q} \mathcal{R}(\alpha, \tau_g, \beta, \gamma, \tau_q)$ is strictly convex in $\alpha \geq 0$. Therefore, its minimizer over $\alpha \geq 0$ is unique. By a similar argument, we show that $\frac{\tau_{g*}}{\beta_*}$ is unique. Consider the change of variable $\tau_g \rightarrow \tilde{\tau}_g = \frac{\tau_g}{\beta}$. Then part of the objective (E.24) that depends on $\tilde{\tau}_g$ can be written as

$$\frac{\beta^2 \tilde{\tau}_g}{2} \frac{n}{d} + \frac{1}{2(\tilde{\tau}_g + 1)} \frac{1}{n} \|\mathbf{w} - \alpha \mathbf{g}\|_{\ell_2}^2 + \frac{1}{n} G_n(\mathbf{w} - \alpha \mathbf{g}; \tilde{\tau}_g, \gamma, \nu) = \frac{\beta^2 \tilde{\tau}_g}{2} \frac{n}{d} + \frac{1}{n} e_f(\mathbf{w} - \alpha \mathbf{g}; \tilde{\tau}_g),$$

using (E.19). As explained above, this converges to the expected Moreau envelope, which is strictly convex in $\tilde{\tau}_g$. Following by the same reasoning for α , one can show that $\min_{\alpha > 0} \max_{0 \leq \beta, \gamma, \tau_q} \mathcal{R}(\alpha, \tilde{\tau}_g, \beta, \gamma, \tau_q)$ is strictly convex in $\tilde{\tau}_g > 0$. Therefore, its minimizer over $\tilde{\tau}_g > 0$ is unique.

APPENDIX F: SOME USEFUL LEMMAS

Here we state some of the technical lemmas that are used in deriving our analytical results.

The first lemma is about the Stieltjes transform of the Marchenko-Pastur distribution.

Definition F.1 The Stieltjes transform $S_\rho(z)$ of a measure of density ρ on a real interval I is the function of the complex variable z defined outside I by the formula

$$S_\rho(z) = \int_I \frac{\rho(t) dt}{z - t} \quad z \in \mathbb{C} \setminus I.$$

Lemma F.2 Suppose that $\mathbf{W} \in \mathbb{R}^{N \times d}$ has rows drawn independently from unit sphere. As $N, d \rightarrow \infty$ and $N/d \rightarrow \psi_1$, the spectral density of $\mathbf{W} \mathbf{W}^\top$ converges (in weak topology in distribution) to the Marchenko-Pastur distribution with Stieltjes transform given by

$$S(z; \psi_1) = \frac{1 - \psi_1 - z - \sqrt{(1 - \psi_1 - z)^2 - 4\psi_1 z}}{-2\psi_1 z},$$

for $z < 0$.

Proof We refer to [2, page 52] for the proof of this proposition. ■

The next lemma is about the spectrum of matrix $\mathbf{J}$ given by

$$(F.1) \quad \mathbf{J} = \left(\mathbf{W}\mathbf{W}^\top \odot \left(\frac{\pi - \cos^{-1}(\mathbf{W}\mathbf{W}^\top)}{2\pi} \right) \right)^{1/2}.$$

Lemma F.3 *Suppose that $\mathbf{W} \in \mathbb{R}^{N \times d}$ has rows chosen randomly and independently of data from the unit sphere, $\text{Unif}(\mathbb{S}^{d-1})$. Let $\mathbf{J}$ be given by (F.1) and suppose that $N/d \rightarrow \psi_1 \in (0, \infty)$, as $n \rightarrow \infty$. Then, the matrix $\mathbf{J}^2$ can (in probability) be approximated consistently in operator norm by the matrix $\mathbf{K}$ given by*

$$\mathbf{K} = \frac{1}{4}(\mathbf{W}\mathbf{W}^\top + \mathbf{I}).$$

In other words, $\|\mathbf{J}^2 - \mathbf{K}\| \rightarrow 0$, in probability, when $n \rightarrow \infty$.

Proof The claim follows from the result of [26, Theorem 2.1] about the spectrum of inner product kernel random matrices, specialized to matrix $\mathbf{J}^2$. Specifically, let $f(z) = z(\pi - \cos^{-1}(z))/(2\pi)$. Then $\mathbf{J}_{ij}^2 = f(\mathbf{w}_i^\top \mathbf{w}_j)$. By employing [26, Theorem 2.1], the kernel matrix $\mathbf{J}^2$ can (in probability) be approximated consistently in operator norm by the matrix $\mathbf{K}$, given by

$$\mathbf{K} = f(0)\mathbf{1}\mathbf{1}^\top + f'(0)\mathbf{W}\mathbf{W}^\top + (f(1) - f(0) - f'(0))\mathbf{I}.$$

For our specific f we have $f(0) = 0$, $f(1) = 1/2$, $f'(0) = 1/4$. ■