
A Precise Performance Analysis of Support Vector Regression

Housseem Sifaou¹ Abba Kammoun¹ Mohamed-Slim Alouini¹

Abstract

In this paper, we study the hard and soft support vector regression techniques applied to a set of n linear measurements of the form $y_i = \beta_\star^T \mathbf{x}_i + n_i$ where $\beta_\star$ is an unknown vector, $\{\mathbf{x}_i\}_{i=1}^n$ are the feature vectors and $\{n_i\}_{i=1}^n$ model the noise. Particularly, under some plausible assumptions on the statistical distribution of the data, we characterize the feasibility condition for the hard support vector regression in the regime of high dimensions and, when feasible, derive an asymptotic approximation for its risk. Similarly, we study the test risk for the soft support vector regression as a function of its parameters. Our results are then used to optimally tune the parameters intervening in the design of hard and soft support vector regression algorithms. Based on our analysis, we illustrate that adding more samples may be harmful to the test performance of support vector regression, while it is always beneficial when the parameters are optimally selected. Such a result reminds a similar phenomenon observed in modern learning architectures according to which optimally tuned architectures present a decreasing test performance curve with respect to the number of samples.

1. Introduction

Motivation. Recent works have demonstrated that the test performance of modern learning architectures exhibits both model-wise and sample-wise double descent phenomena that defy conventional statistical intuition. Model-wise descent, reported in recent works (Belkin et al., 2019a;c; Geiger et al., 2019), suggests that, for very large architectures, performance improves with the number of parameters, thus contradicting the bias-variance trade-off. On the other hand, sample-wise descent, discussed recently in the works

of Nakkiran *et al.* (Nakkiran et al., 2020a;b) indicates that more data may harm the performance. One potential solution to avoid such a behavior consists in optimally tuning the involved parameters. In doing so, the test performance in most scenarios decreases with the number of samples.

In this paper, we investigate the sample-wise double descent phenomenon for basic linear models. More precisely, we assume independent data samples $(\mathbf{x}_i, y_i)$, $i = 1, \dots, n$ distributed as:

$$y_i = \beta_\star^T \mathbf{x}_i + \sigma n_i,$$

where $\mathbf{x}_i \in \mathbb{R}^p$ is the feature vector assumed to have zero mean and covariance $\mathbf{I}_p$, y_i represent the scalar response variables while σn_i stands for zero-mean noise with variance σ^2 . To estimate $\beta_\star$, we consider support vector regression techniques: hard-support vector regression (H-SVR), which estimates the regression vector β with minimum ℓ_2 norm that satisfies the constraints $y_i = \beta^T \mathbf{x}_i$ up to a maximum error ϵ , and soft-support regression (S-SVR) which uses a regularization constant C that aims to create a trade-off between the minimization of the training error and the minimization of the model complexity. SVR has been extensively used in several applications that vary from biomedical analysis (Hamdi et al., 2018) to financial time series (Ju et al., 2014; Yang et al., 2009; Qu & Zhang) and weather forecasting (Guajardo et al., 2006).

We study the performance of H-SVR and S-SVR when the number of features p and the sample size n grow simultaneously large such that $\frac{n}{p} \rightarrow \delta$ with $\delta > 0$ and the norm of $\|\beta_\star\|$ converges to β . One major outcome of the present work is to recover interesting behaviors observed in large-scale machine learning architectures. Particularly, we show numerically that the double descent behavior appears only when the H-SVR or S-SVR parameters are not properly tuned. Such a behavior reminds the recent findings in (Nakkiran et al., 2020b) that suggest that unregularized models often suffer from the sample-wise double descent phenomenon, while optimally tuned models usually present a monotonic risk with respect to the number of samples.

Contributions. This paper investigates the test risk behavior as a function of the sample size for H-SVR and S-SVR techniques. Contrary to least squares regression, which involves explicit form expressions for the solution, H-SVR and S-SVR require solving convex-optimization problems,

¹Computer, Electrical, and Mathematical Sciences & Engineering Division, King Abdullah University of Science and Technology (KAUST), Thuwal, Saudi Arabia. Correspondence to: Housseem Sifaou <housssem.sifaou@kaust.edu.sa>.

which do not have closed-form solutions. To analyze the test risk, we rely on the Gaussian min-max theorem (CGMT) framework (Thrampoulidis et al., 2018) and more specifically on the extension of this framework recently developed in (Deng et al., 2020), which has been proven to be suitable to analyze functionals of solutions of convex optimizations problems.

Concretely, our results for the H-SVR and S-SVR can be summarized as follows:

1. We derive for a fixed error tolerance ϵ , a sharp phase transition-threshold δ_* beyond which the H-SVR becomes infeasible. Interestingly, we illustrate that the transition threshold depends only on ϵ and the noise variance and not on the SNR defined as $\text{SNR} := \frac{\beta^2}{\sigma^2}$. Moreover, we show that δ_* is always greater than 1, which should be compared with the condition $\delta < 1$ required for the least square estimator to exist. As a side note, we prove that contrary to hard-margin support vector classifiers, H-SVR can always be feasible through a proper tuning of the tolerance error ϵ . This allows us to study the test risk of the H-SVR as a function of δ when $\delta \in (0, \infty)$ and ϵ carefully tuned to satisfy the feasibility condition.
2. For fixed error tolerance ϵ , we numerically illustrate that for moderate to large SNR the test risk of the H-SVR is a non-monotonic curve presenting a unique minimum that becomes the closest to δ_* as the SNR increases. For low SNR values, the test risk is an increasing function of δ_* and is always worse than the null risk associated with the null estimator. It is worth mentioning that behavior of the same kind was reported for the min-norm least square estimator in (Hastie et al., 2019). Additionally, we illustrate that when the parameter ϵ is optimally tuned, the test risk becomes a decreasing function of δ and equivalently of the number of data samples.
3. Similarly, we derive the expression for the asymptotic test risk as a function of ϵ , δ , and the regularization constant C . Without optimal tuning of the regularization constant C and factor ϵ , the test curve as a function of the sample test size presents a double descent, which disappears when optimal settings of these constants is considered.
4. We study the robustness of the S-SVR and H-SVR to impulsive noise. We illustrate that contrary to H-SVR, S-SVR, when optimally tuned, is resilient to impulsive noises. Particularly, we show that for mild impulsive noise conditions, S-SVR presents a slightly lower risk than optimally tuned ridge regression estimators but largely outperforms it under moderate to severe impulsive noise conditions.

Related works. The present work is part of the continued efforts to understand the double descent phenomena in large-scale machine learning architectures. An important body of research works focused on establishing the behavior of double descent of the test risk as a function of the model size in a variety of machine learning algorithms (Belkin et al., 2019a; Bös & Opper, 1997; Spigler et al., 2019). Very recently, the work in (Nakkiran et al., 2020a) discovered that double descent occurs not just as a function of the model size but also as a function of the sample size (Nakkiran et al., 2020a). A major consequence of such a behavior is that performance may be degraded as we increase the number of samples.

To further understand the generalization error, several works considered to analyze it as a function of the model size for mathematically tractable settings in regression (Hastie et al., 2019; Belkin et al., 2019b; Muthukumar et al., 2020; Mitra, 2019; Candes & Sur, 2018) and more recently in classification (Deng et al., 2020; Kini & Thrampoulidis, 2020; Montanari et al., 2020), with the goal of investigating under which conditions, the double descent occurs. In this paper, similarly to (Nakkiran et al., 2020b), we instead focus on the effect of the sample size on the test performance, but with the H-SVR and S-SVR as case examples. Moreover, on the technical level, our analysis provides sharp characterizations of the performance using the recently developed extension of the CGMT framework (Deng et al., 2020).

2. Problem formulation

Consider the problem of estimating the scalar response y of a vector $\mathbf{x}$ in $\mathbb{R}^p$ from a set of n data samples $\{(\mathbf{x}_i, y_i)\}_{i=1}^n$ following the linear model:

$$y_i = \beta_*^T \mathbf{x}_i + \sigma n_i \quad (1)$$

where β_* is an unknown vector, $\{\sigma n_i\}_{i=1}^n$ represent noise samples with zero-mean and variance σ^2 . We further assume that $\mathbf{x}_i$ is Gaussian with zero mean and covariance $\mathbf{I}_p$. To estimate β_* , we consider support vector regression methods, namely the hard support vector regression denoted by H-SVR and the soft support vector regression referred to as S-SVR. The H-SVR looks for a function $y = \beta^T \mathbf{x}$ such that all data points $(\mathbf{x}_i, \beta^T \mathbf{x}_i)$ deviates at most ϵ from their targets y_i . Formally, this regression problem can be written as:

$$\begin{aligned} \hat{\mathbf{w}}_H &:= \arg \min_{\mathbf{w}} \frac{1}{2} \|\mathbf{w}\|^2 \\ \text{s.t. } & |y_i - \mathbf{w}^T \mathbf{x}_i| \leq \epsilon \end{aligned} \quad (2)$$

where $\|\cdot\|$ denotes the ℓ_2 -norm of a vector. It is worth mentioning that when $\epsilon = 0$ and $n \leq p$, the H-SVR boils down to the least square estimator. In this case, it perfectly interpolates the training data, satisfying $y_i = \mathbf{x}_i^T \hat{\mathbf{w}}_H$, $i = 1, \dots, n$. In general, depending on the value of ϵ , there

may not be a solution that satisfies the constraints. As in support vector machines for classification, one solution to deal with such cases is to add slack variables that while relaxing the constraints, penalize, in the objective function, large deviations from them. Applying this approach gives the S-SVR method which involves solving the following optimization problem:

$$\begin{aligned} \hat{\mathbf{w}}_S := \arg \min_{\mathbf{w}} \quad & \frac{1}{2} \|\mathbf{w}\|^2 + \frac{C}{p} \sum_{i=1}^n (\xi_i + \tilde{\xi}_i) \\ \text{s.t.} \quad & y_i - \mathbf{w}^T \mathbf{x}_i \leq \epsilon + \xi_i, \quad i = 1, \dots, n \quad (3) \\ & \mathbf{w}^T \mathbf{x}_i - y_i \leq \epsilon + \tilde{\xi}_i \\ & \xi_i, \tilde{\xi}_i \geq 0. \end{aligned}$$

The aim of the present work is to characterize analytically the performance of the H-SVR and the S-SVR. The assumption underlying the analysis is to consider that the number of samples and that of features grow with the same pace, and will be made more specific in the sequel.

Prediction risk. The metric of interest in this paper is the prediction risk. For a given estimator $\hat{\beta}$, the prediction risk is defined as:

$$\mathcal{R}(\hat{\beta}) := \mathbb{E}_{\mathbf{x}, y} |\mathbf{x}^T \hat{\beta} - \mathbf{x}^T \beta_\star|^2 = \|\hat{\beta} - \beta_\star\|^2,$$

where $\mathbf{x}$ and y are test points following the model (1) but are independent of the training set. Expressing $\mathcal{R}(\hat{\beta})$ as:

$$\mathcal{R}(\hat{\beta}) = \|\beta_\star\|^2 + \|\hat{\beta}\|^2 - 2\|\beta_\star\| \|\hat{\beta}\| \frac{\beta_\star^T \hat{\beta}}{\|\beta_\star\| \|\hat{\beta}\|},$$

we can easily see that the risk depends on $\hat{\beta}$ through its norm $\|\hat{\beta}\|$ and the cosine similarity between $\beta_\star$ and $\hat{\beta}$:

$$\cos(\beta_\star, \hat{\beta}) := \frac{\beta_\star^T \hat{\beta}}{\|\beta_\star\| \|\hat{\beta}\|}.$$

3. Main results

In this paper, we consider studying the performance of the hard and soft support vector regression problems under the asymptotic regime in which n and p grow large at the same pace. More specifically, the following assumption is considered.

Assumption 1. *Our study is based on the following set of assumptions:*

- n and p grow to infinity with $\frac{n}{p} \rightarrow \delta$.
- The noise variance σ^2 is a fixed positive constant.
- $\|\beta_\star\| \rightarrow \beta$ where β is a certain positive scalar.

- **Data model:** The training $\{\mathbf{x}_i\}_{i=1}^n$ are independent and identically distributed following standard normal distribution. Moreover, $\{n_i\}$ are independent and drawn from a zero-mean unit-variance symmetric distribution p_N .

3.1. Hard SVR

As mentioned earlier, the H-SVR problem is not always feasible. We provide in Theorem 1 a sharp characterization of the feasibility region of the H-SVR in the asymptotic regime defined in Assumption 1.

Theorem 1 (Feasibility of the H-SVR). *Let $\delta_\star$ be defined as:*

$$\delta_\star = \frac{1}{\inf_{t \in \mathbb{R}} \mathbb{E} \left(|G + tN| - t \frac{\epsilon}{\sigma} \right)_+^2}, \quad (4)$$

where $(x)_+ \triangleq \max(x, 0)$ and the expectation is taken over the distribution of G and N where $G \sim \mathcal{N}(0, 1)$ and $N \sim p_N$. Consider the asymptotic regime and data model in Assumption 1. Then, the following statements hold true:

$$\delta > \delta_\star \Rightarrow \mathbb{P}[\text{The H-SVR is feasible for suff. large } n] = 0 \quad (5)$$

$$\delta < \delta_\star \Rightarrow \mathbb{P}[\text{The H-SVR is feasible for suff. large } n] = 1 \quad (6)$$

The proofs can be found in the supplementary material.

Remark 1. (Feasibility of the H-SVR depends on the noise variance but not on the SNR.) The above result establishes that the existence of the H-SVR undergoes a sharp transition phenomenon. Particularly, in the limit of large sample size n and number of features p such that $\frac{n}{p} \rightarrow \delta$, the H-SVR is almost surely unfeasible when $\delta > \delta_\star$ and always feasible when $\delta < \delta_\star$. The obtained expression is reminiscent of other previously established result for the existence of the hard-margin SVM for classification established in a series of recent works (Kammoun & Alouini, 2020b) and (Deng et al., 2020). However, contrary to the expressions obtained in these works, the separability boundary curve captured by $\delta_\star$ does not depend on the Euclidean norm of $\beta_\star$, or equivalently on the SNR defined as $\frac{\|\beta_\star\|^2}{\sigma^2}$ but only on the noise variance. The reason behind this is that feasibility is essentially related to how much the data samples deviate from the hyperplane defined as $\beta_\star^T \mathbf{x} = y$. We note that as σ approaches 0 and $\epsilon \neq 0$, $\delta_\star \rightarrow \infty$, which implies that the H-SVR is always feasible in this case. This is because in the noiseless case, all data samples $(\mathbf{x}_i, y_i)$ belong to the hyperplane $y_i = \beta_\star^T \mathbf{x}_i$ and thus $\beta_\star$ is in the feasibility set of the H-SVR regardless of the value of ϵ and also on β . On the other hand, as σ increases, $\delta_\star$ decreases, which suggests that the H-SVR becomes less feasible since it is more difficult to find a hyperplane that contains all data samples with a reasonable error tolerance ϵ .

Remark 2. (The H-SVR can be feasible when the least square estimator is not) One can easily check that if $\epsilon = 0$, then $\delta_* = 1$. This result makes sense since as long as $\delta < 1$, the linear system $\beta^T \mathbf{x}_i = y_i$ is under determined and as such a solution β exists. However, when $\delta_* > 1$, the linear system becomes over-determined, and as such it is impossible to find a solution to this linear system. Moreover, since δ_* is an increasing function of ϵ , we conclude that $\delta_* > 1$ for all $\epsilon > 0$. This particularly shows that the H-SVR provides a larger feasibility region than the least square estimator, for which δ should be less than 1 to exist. We can even push this result further and claim that *every δ is feasible once ϵ is appropriately tuned*. To see this, it suffices to note that $\epsilon \mapsto \delta_*$ is an increasing function establishing a one-to-one map from $(0, \infty)$ to $(1, \infty)$. Hence, for any $\delta \in (1, \infty)$ there exists $\epsilon^*(\delta)$ such that for all $\epsilon > \epsilon^*(\delta)$, the H-SVR is almost surely feasible.

For the sake of illustration, Figure 1 displays δ_* as a function of ϵ for several values of σ . As can be seen, δ_* is an increasing function growing to infinity with ϵ . The value of ϵ plays a fundamental role to remediate the effect of the noise and ensure the feasibility of the H-SVR. One can note as expected that δ_* for $\epsilon = 0.1$ and $\sigma = 0.1$ is the same as the one obtained when $\epsilon = 0.2$ and $\sigma = 0.2$. This finding can be easily concluded from (4). Moreover, in agreement with our previous discussion, we can easily see that δ_* decreases significantly as the noise variance increases. Such a scenario can be fixed by adapting the value of ϵ .

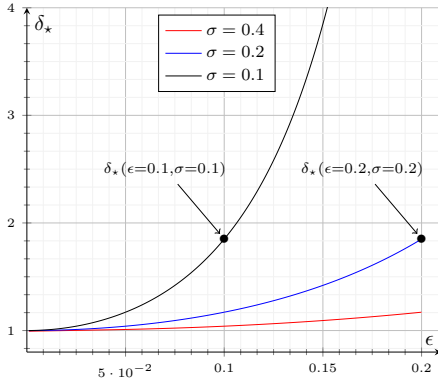


Figure 1. Theoretical predictions of δ_* as a function of ϵ for different noise variance values. The figure shows that every δ can be forced to be in the feasibility region by appropriately choosing ϵ .

Having characterized the feasibility region of the H-SVR, we are now ready to provide sharp asymptotics of its performance in terms of the prediction risk and the cosine similarity.

Theorem 2 (Asymptotic prediction risk of H-SVR). *Define*

function $\mathcal{D} : \mathbb{R}^2 \rightarrow \mathbb{R}$ as:

$$\mathcal{D}(\tilde{\gamma}_1, \tilde{\gamma}_2) = \sqrt{\delta} \sqrt{\mathbb{E} \left(\left| \sqrt{\tilde{\gamma}_1^2 + \tilde{\gamma}_2^2} G + N \right| - \frac{\epsilon}{\sigma} \right)_+^2} - \tilde{\gamma}_1$$

where the expectation is taken over the distribution of the independent random variables G and N drawn respectively from the standard normal distribution $\mathcal{N}(0, 1)$ and p_N . Let $\tilde{\gamma}_1^*$ and $\tilde{\gamma}_2^*$ be the unique solutions to the following optimization problem:

$$(\tilde{\gamma}_1^*, \tilde{\gamma}_2^*) = \arg \min_{\substack{\tilde{\gamma}_1, \tilde{\gamma}_2 \\ \mathcal{D}(\tilde{\gamma}_1, \tilde{\gamma}_2) \leq 0}} \frac{1}{2} \left(\tilde{\gamma}_2 - \frac{\beta}{\sigma} \right)^2 + \frac{1}{2} \tilde{\gamma}_1^2. \quad (7)$$

Let $\hat{\mathbf{w}}_H$ be the solution to the H-SVR. Then, under Assumption 1 and assuming $\delta < \delta_*$, the following convergences hold true:

$$\mathcal{R}(\hat{\mathbf{w}}_H) - \bar{R}_H \xrightarrow{a.s.} 0$$

where $\bar{R}_H = \sigma^2((\tilde{\gamma}_1^*)^2 + (\tilde{\gamma}_2^*)^2)$ and

$$\frac{\hat{\mathbf{w}}_H^T \beta_*}{\|\hat{\mathbf{w}}_H\| \|\beta_*\|} \xrightarrow{a.s.} \frac{\frac{\beta}{\sigma} - \tilde{\gamma}_2^*}{\sqrt{(\tilde{\gamma}_1^*)^2 + (\tilde{\gamma}_2^* - \frac{\beta}{\sigma})^2}}.$$

Remark 3. (Behavior of the H-SVR as $\delta \rightarrow 0$). As δ tending to zero, one can check after careful investigation of the asymptotic expressions that $\tilde{\gamma}_2^* \rightarrow \frac{\beta}{\sigma}$ and $\tilde{\gamma}_1^* \rightarrow 0$. In this case, the asymptotic risk is thus given by β^2 . To see this, it suffices to note that $\mathcal{D}(\delta^{\frac{1}{4}}, \frac{\beta}{\sigma})$ converges from below to zero as $\delta \downarrow 0$. By continuity of $\mathcal{D}$, we may find η sufficiently small such that, for all $(\tilde{\gamma}_1, \tilde{\gamma}_2) \in \mathcal{C}(\eta)$ where

$$\mathcal{C}(\eta) = \left\{ (\tilde{\gamma}_1, \tilde{\gamma}_2) \mid \tilde{\gamma}_1 \in (0, \eta), \tilde{\gamma}_2 \in \left(\frac{\beta}{\sigma} - \eta, \frac{\beta}{\sigma} + \eta \right) \right\},$$

we have $\mathcal{D}(\tilde{\gamma}_1, \tilde{\gamma}_2) \leq 0$. Moreover, it is easy to see that the objective in (7) can be bounded by η^2 when $(\tilde{\gamma}_1, \tilde{\gamma}_2) \in \mathcal{C}(\eta)$. Evaluation of this objective when $|\tilde{\gamma}_2 - \frac{\beta}{\sigma}| \geq \sqrt{2}\eta$ or when $|\tilde{\gamma}_1| \geq \sqrt{2}\eta$ yields values greater than η^2 . Hence, necessarily, $(\tilde{\gamma}_1, \tilde{\gamma}_2) \in \mathcal{C}(\sqrt{2}\eta)$, which proves the desired. Finally, observing that the risk of the null estimator $\hat{\beta} = 0$ is also β^2 , we conclude that the H-SVR is worse than the null estimator when a small number of samples is employed.

Remark 4. (Behavior of the H-SVR as $\delta \rightarrow \delta_*$). As $\delta \rightarrow \delta_*$, the set $\{(\tilde{\gamma}_1, \tilde{\gamma}_2) \mid \mathcal{D}(\tilde{\gamma}_1, \tilde{\gamma}_2) \leq 0\}$ becomes the unit set $\{(\tilde{\gamma}_1^*, 0)\}$ where $\tilde{\gamma}_1^*$ is the smallest solution to the equation $\mathcal{D}(\tilde{\gamma}_1, 0) = 0$. The asymptotic risk thus becomes equal to $\sigma^2(\tilde{\gamma}_1^*)^2$ and is as such independent of the SNR $\frac{\beta^2}{\sigma^2}$. As a result, when δ approaches δ_* , it is the noise variance and the value of ϵ that determines the performance of the H-SVR and not the SNR. This is in opposition to the behavior in the operation region $\delta \rightarrow 0$, for which the risk tends to β^2 .

3.2. Soft SVR

In this section, the soft SVR problem is considered and the convergence of the corresponding prediction risk is established.

Theorem 3 (Asymptotic prediction risk of S-SVR). *Under Assumption 1, we have*

$$\mathcal{R}(\hat{\mathbf{w}}_S) \xrightarrow{a.s.} \sigma^2(\tilde{\gamma}_1^{*2} + \tilde{\gamma}_2^{*2})$$

where $\tilde{\gamma}_1^*$ and $\tilde{\gamma}_2^*$ are the solutions of the following scalar optimization problem

$$\bar{\phi} = \min_{\tilde{\gamma}_1, \tilde{\gamma}_2} \sup_{\chi > 0} \bar{D}(\tilde{\gamma}_1, \tilde{\gamma}_2, \chi)$$

with

$$\begin{aligned} \bar{D}(\tilde{\gamma}_1, \tilde{\gamma}_2, \chi) &= \frac{1}{2}\tilde{\gamma}_1^2 + \frac{1}{2}\left(\tilde{\gamma}_2 - \frac{\beta}{\sigma}\right)^2 - \frac{\tilde{\gamma}_1\chi}{2\sigma} \\ &+ \frac{\delta}{\sigma}\mathbb{E}\left\{C\left[X - \frac{C\tilde{\gamma}_1}{2\chi}\right]\mathbb{1}_{\{X\chi > \tilde{\gamma}_1 C\}} + \frac{\chi}{2\tilde{\gamma}_1}X^2\mathbb{1}_{\{X\chi \leq \tilde{\gamma}_1 C\}}\right\} \end{aligned}$$

where $X = \left(\left|\sqrt{\tilde{\gamma}_1^2 + \tilde{\gamma}_2^2}G + N\right| - \epsilon/\sigma\right)_+$ and the expectation is with respect to the distributions of G and N with $G \sim \mathcal{N}(0, 1)$ and $N \sim p_N$.

Remark 5. Theorem 3 can be used to optimally tune the parameters (ϵ, C) so that they minimize the asymptotic test risk. For that, one is required to estimate the noise variance σ^2 and the signal power β^2 . These can be easily estimated through the following approach. Let $\mathbf{y} = [y_1, \dots, y_n]^T$ be the vector of the responses and $\mathbf{X} = [\mathbf{x}_1, \dots, \mathbf{x}_n]$ the matrix stacking all training samples. Then,

$$\mathbf{y} = \mathbf{X}^T \boldsymbol{\beta}_* + \sigma \mathbf{n},$$

where $\mathbf{n} = [n_1, \dots, n_n]^T$. From the strong law of large numbers,

$$\frac{1}{n}\mathbf{y}^T \mathbf{y} \xrightarrow{a.s.} \sigma^2 + \beta^2. \quad (8)$$

On the other hand, assuming $\delta > 1$,

$$\frac{1}{n}\mathbf{y}^T (\mathbf{I}_n - \mathbf{X}^T (\mathbf{X}\mathbf{X}^T)^{-1} \mathbf{X}) \mathbf{y} \xrightarrow{a.s.} \sigma^2 \left(1 - \frac{1}{\delta}\right). \quad (9)$$

Combining (8) and (9), consistent estimators for the noise variance σ^2 and for β^2 are given by:

$$\hat{\sigma}^2 = \frac{\frac{1}{n}\mathbf{y}^T (\mathbf{I}_n - \mathbf{X}^T (\mathbf{X}\mathbf{X}^T)^{-1} \mathbf{X}) \mathbf{y}}{1 - \frac{1}{\delta}}, \quad (10)$$

$$\hat{\beta}^2 = \frac{1}{n}\mathbf{y}^T \mathbf{y} - \frac{\frac{1}{n}\mathbf{y}^T (\mathbf{I}_n - \mathbf{X}^T (\mathbf{X}\mathbf{X}^T)^{-1} \mathbf{X}) \mathbf{y}}{1 - \frac{1}{\delta}}. \quad (11)$$

In case $\delta < 1$, the problem of estimating the noise variance becomes more challenging, and there are, to the best of our knowledge, no general unbiased estimators with the same statistical guarantees as in the case $\delta > 1$. To address this issue, some other techniques may be used (Cherkassky & Ma, 2004) but they are not guaranteed to lead to consistent estimators.

Remark 6. Behavior of the S-SVR when $\delta \rightarrow 0$. The test risk of the S-SVR is much more involved than that of the H-SVR. Nevertheless, it can be easily seen that when δ goes to zero,

$$\limsup_{\delta \rightarrow 0} \sup_{\chi \geq 0} \bar{D}(\tilde{\gamma}_1, \tilde{\gamma}_2, \chi) \stackrel{(a)}{=} \inf_{\delta \geq 0} \sup_{\chi \geq 0} \bar{D}(\tilde{\gamma}_1, \tilde{\gamma}_2, \chi) \quad (12)$$

$$\stackrel{(b)}{=} \sup_{\chi \geq 0} \inf_{\delta \geq 0} \bar{D}(\tilde{\gamma}_1, \tilde{\gamma}_2, \chi) \quad (13)$$

$$= \sup_{\chi \geq 0} \lim_{\delta \rightarrow 0} \bar{D}(\tilde{\gamma}_1, \tilde{\gamma}_2, \chi) \quad (14)$$

$$= \frac{1}{2}\tilde{\gamma}_1^2 + \frac{1}{2}\left(\tilde{\gamma}_2 - \frac{\beta}{\sigma}\right)^2 \quad (15)$$

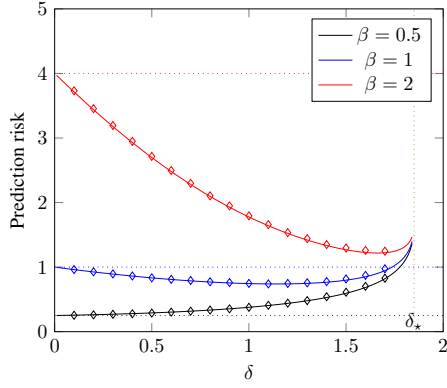
where (a) follows from the fact that the objective function is an increasing function in δ and (b) from the fact that the objective function is convex in δ and concave in χ . The asymptotic limit of $\bar{D}$ in (15) has a unique minimum given by $\tilde{\gamma}_2 = \frac{\beta}{\sigma}$ and $\tilde{\gamma}_1 = 0$. Plugging these values into that of the test risk, we conclude that when δ goes to zero, the test risk of the S-SVR converges to that of the null estimator.

4. Numerical illustration

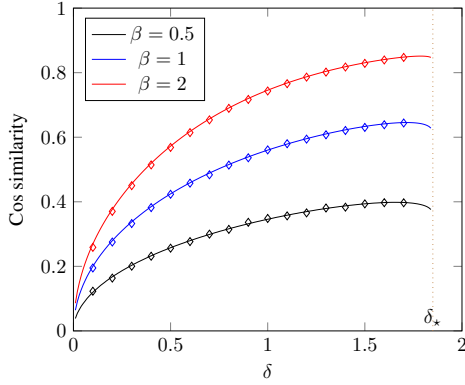
4.1. H-SVR

4.1.1. TEST RISK AS A FUNCTION OF THE NUMBER OF SAMPLES

In a first experiment, we investigate the behavior of the test risk of H-SVR as a function of the number of samples for different values of the signal power $\beta^2 = \|\boldsymbol{\beta}_*\|^2$. Particularly, for each $\beta \in \{0.5, 1, 2\}$, we fix the noise variance σ^2 and ϵ and plot the test risk and cosine similarity over the range $[0, \delta_*]$ over which the H-SVR is feasible. Fig. 2 represents the theoretical results along with their empirical averages obtained for $p = 200$ and $n = \lfloor \delta p \rfloor$. As can be seen, this figure's results validate the accuracy of the theoretical predictions: a perfect match is noted over the whole range of δ and for all signal power values. We also corroborate our predictions in Remark 3 and Remark 4: the risk tends to β^2 which is the null estimator's risk when $\delta \rightarrow 0$, while it tends to the same limit irrespective of β^2 as $\delta \rightarrow \delta_*$. Away from these limiting cases, we note that for moderate to high signal powers, the test risk presents a non-monotonic behavior with respect to δ and as such with respect to the number of samples. The minimal risk corresponds to a δ that becomes the nearest to δ_* as the signal power β^2 increases. However, for low signal powers, the test risk is an increasing function of the number of samples and is always larger than β^2 , which is the null estimator's risk. Such behavior is similar to that of the min-norm least square estimator, which becomes worse than the null estimator when the SNR is less than 1 (Hastie et al., 2019).



(a) Prediction risk



(b) Cos similarity

Figure 2. Performance of H-SVR as a function of δ when $\sigma = 1$, $\epsilon = 1$ for different values of β . The continuous line curves correspond to the asymptotic performance while the points denote finite-sample performance when $p = 200$ and $n = \lfloor \delta p \rfloor$. The null risks (corresponding to $\hat{\mathbf{w}}_H = \mathbf{0}$) are also reported by the dotted lines for the different values of β .

4.1.2. IMPACT OF CHOICE OF ϵ ON THE TEST PERFORMANCE

Fig. 3 displays the test risk with respect to ϵ for fixed signal power and noise variance and oversampling ratios $\delta = 1$ and $\delta = 1.4$, respectively. As can be noted, the test performance is sensitive to the choice of ϵ . An arbitrary choice of ϵ may lead to a significant loss in test performance. Indeed, a small ϵ tolerates less deviation from the plane $y = \hat{\mathbf{w}}_H^T \mathbf{x}$, which becomes inappropriate when the noise variance increases. On the other hand, a larger ϵ tolerates more deviation, and as such, tends to give less credit on the information from the training samples. We can also note that the optimal ϵ increases when more training samples are used. This can be explained by the fact that when using more training samples, it becomes harder to fit them into the insensitivity tube. Moreover, as can be seen from this figure, a right choice for the value ϵ is essential in practice, as arbitrary

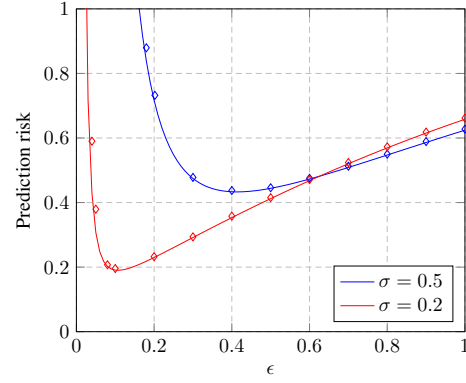
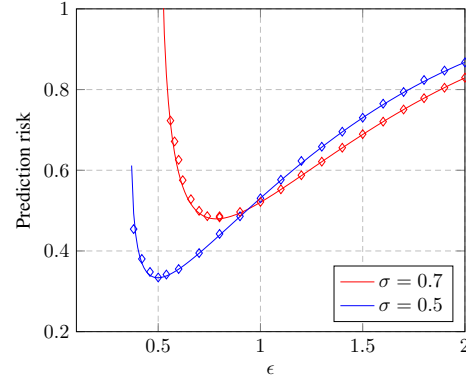

(a) $\delta = 1$

(b) $\delta = 1.4$

Figure 3. Performance of H-SVR vs ϵ when $\beta = 1$ for different values of the noise variance σ^2 and for different values of δ . The continuous line curves correspond to the theoretical predictions while the points denote finite-sample performance when $p = 200$ and $n = \lfloor \delta p \rfloor$.

choices may lead to severe risk performance degradation. Several previous works addressed this question (Cherkassky & Ma, 2004; 2002). However, they do not rely on theoretical analysis but instead on cross-validation approaches. Finally, with reference to **Remark 2**, we plot in Fig. 4 the risk of H-SVR with respect to δ when at each δ , the optimal ϵ that minimizes the optimal risk is used. The obtained results show that the test risk becomes, in this case, a decreasing function of δ . This is in agreement with the fact that in optimally regularized learning architectures, there is always a gain from using more training samples.

4.2. S-SVR

4.2.1. IMPACT OF THE PARAMETERS ϵ AND C

In Fig. 5 and Fig. 6, we investigate the effect of the hyperparameters C and ϵ on the performance of soft SVR. As shown in these figures, arbitrary choices for the pair (ϵ, C) may lead to a significant degradation in the test risk per-

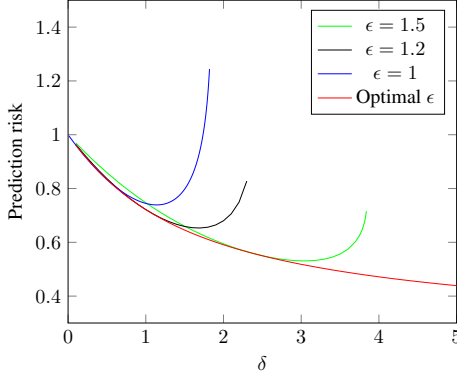


Figure 4. Performance of H-SVR as a function of δ when $\sigma = 1$ and $\beta = 1$ for different values of ϵ . All the curves correspond to the asymptotic performance predictions.

formance compared to the optimal performance associated with optimal selection of (ϵ, C) . Thus, our results again emphasize the importance of the appropriate selection of these parameters and suggest the practical relevance of theoretical aided approaches to select these parameters that may complement existing cross-validation techniques.

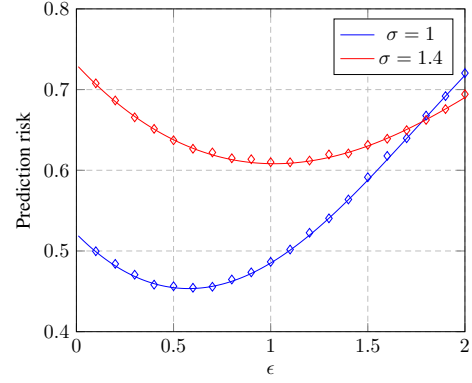
4.2.2. SAMPLE-WISE DESCENT PHENOMENON

In Fig. 7, we plot the test risk with respect to δ for the S-SVR for different choices of C and fixed ϵ . As can be seen, for δ tending to zero, the S-SVR test risk converges to that of the null risk, which was also predicted by our analysis in **Remark 6**. Moreover, depending on the value of C , the test risk may manifest a cusp at $\delta \sim \delta_*$ that becomes more pronounced as C increases. This can be explained by the fact that as C increases, the behavior of S-SVR approaches that of H-SVR, for which the problem becomes unfeasible when $\delta > \delta_*$. We also note that the choice of the parameter C plays a fundamental role in the test risk performance. Optimal values of C always guarantee that the test risk performance decreases with more training samples being used, while arbitrary choices can lead to the test risk increasing for more training samples. This emphasizes the importance of the appropriate selection of the parameter C to avoid the double descent phenomenon.

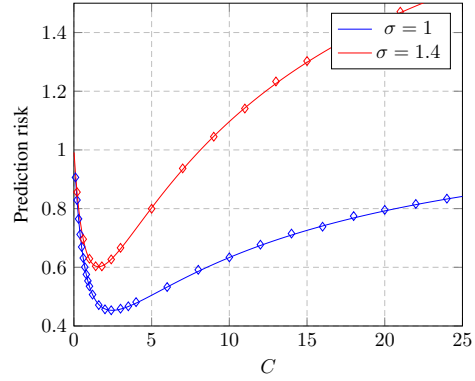
4.3. Comparison with ridge regression estimators under impulsive noises:

In this experiment, we investigate S-SVR and H-SVR's resilience when optimally designed (optimal ϵ for H-SVR and optimal C and ϵ for S-SVR) to impulsive noises and compare them to the ridge regression with optimal regularization. Particularly, we consider the case in which the noise is sampled from the distributional model:

$$n = \sqrt{\tau} \mathcal{N}(0, 1)$$



(a) vs. ϵ for $C = 2.4$



(b) vs. C for $\epsilon = 0.6$

Figure 5. Prediction risk of S-SVR vs ϵ and C when $\delta = 2$, $\beta = 1$ and $\sigma = 1$. The continuous line curves correspond to the theoretical predictions while the points denote finite-sample performance when $p = 200$ and $n = \lfloor \delta p \rfloor$.

where τ follows an inverse Gamma distribution with shape $\frac{d}{2}$ and scale $\frac{2}{d}$, that is $\tau \sim \frac{d}{\chi_d^2}$ with χ_d^2 being the chi-square distribution with d degrees of freedom.

Fig. 8 represents the test risk performance of the aforementioned estimators for $d = 3$ and $d = 10$. As can be seen, H-SVR is very sensitive to impulsive noises with a performance approaching that of the null estimator for highly impulsive noises. This behavior can be explained by the fact that in highly impulsive noises (small d), H-SVR needs a very large ϵ to guarantee that outliers satisfy the feasibility conditions. However, with a large ϵ , the constraints in (2) becomes irrelevant for the remaining well-behaved observations. This favors the H-SVR to select the null estimator as it would minimize the objective in (2) while satisfying the constraints. On the other hand, the S-SVR overcomes such behavior since it does not have to satisfy the constraints of (2). By selecting the parameter C to the value that minimizes the test risk, it will adaptively control the effect of outliers by relaxing the most unlikely constraints in (2). Moreover,

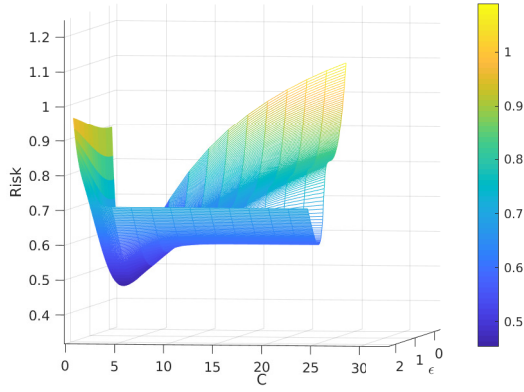
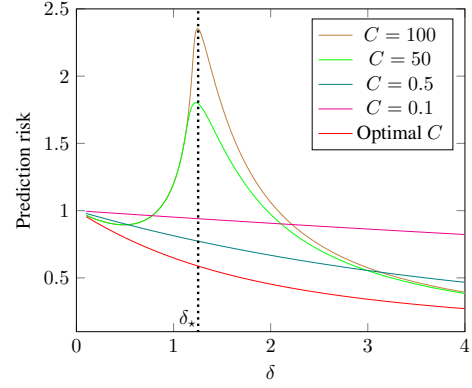


Figure 6. Prediction risk vs (C, ϵ) for $\delta = 2$, $\beta = 1$ and $\sigma = 1$.

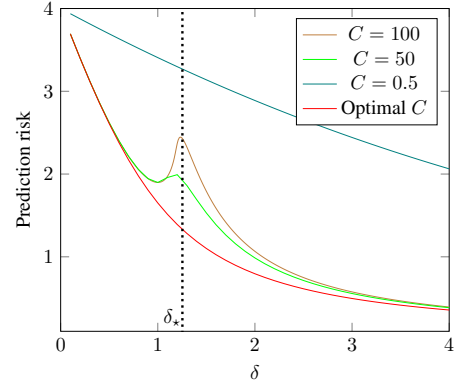
it can be seen that, although S-SVR is slightly less efficient than the ridge regression estimator in mild impulsive noises, it presents a much lower risk under highly impulsive noises. While the robustness of regularized support vector machine methods for both regression and classification is a well-known fact reported in many previous works (Xu et al., 2009; Hable & Christmann, 2011), our result contributes to quantitatively assess such robustness by measuring the test risk under impulsive noises.

5. Conclusion

In this paper, we studied the asymptotic test risk of hard and soft support vector regression techniques with isotropic Gaussian features and under symmetric noise distributions in the regime of high dimensions. We used these results to illustrate the impact of the intervening parameters on the test risk behavior. Particularly, we demonstrate that arbitrary choices of the parameters of the hard SVR and the soft SVR may lead to the test risk presenting a non-monotonic behavior as a function of the number of samples, which illustrates the fact that adding more samples may be harmful to the performance. On the contrary, we show that optimally-tuned hard SVR and soft SVR present a decreasing test risk curve, which shows the importance of carefully selecting their parameters to minimize the test risk and guarantee the positive impact of more data on the test risk performance. Our findings are consistent with similar results obtained for linear regression and neural networks (Nakkiran et al., 2020b). However, as compared to linear regression, we demonstrate that soft support vector regression with optimal regularization is more robust to the presence of outliers, corroborating similar previous findings in earlier works in (Xu et al., 2009; Hable & Christmann, 2011). Several extensions of our work are worth investigating. One important research direction is to understand the effect of correlated features on the test risk of hard and support regression techniques. Some re-



(a) $\beta = 1$



(b) $\beta = 2$

Figure 7. Prediction risk of S-SVR vs. δ for different values of C when $\epsilon = 0.6$ and $\sigma = 1$. Illustration of the sample-wise double decent and how optimal regularization mitigates it.

cent works have investigated the role of correlation and regularization on the test risk for linear regression models (Lolas, 2020; Kobak et al., 2020). A significant advantage of such an analysis is that it can illustrate the importance of investing efforts in theoretically-aided approaches to assist in setting the regularization parameters. Another important research direction is investigating the use of kernel support vector regression methods and understanding their underlying mechanisms to handle involved non-linear data models.

References

- Bai, Z. and Silverstein, J. W. Spectral Analysis of Large Dimensional Random Matrices. *Springer Series in Statistics*, 2009.
- Belkin, M., Hsu, D., Ma, S., and Mandal, S. Reconciling modern machine-learning practice and the classical bias-variance trade-off. *Proceedings of the National Academy of Sciences*, 116(32):15849–15854, 2019a. ISSN 0027-8424. doi: 10.1073/pnas.1903070116.

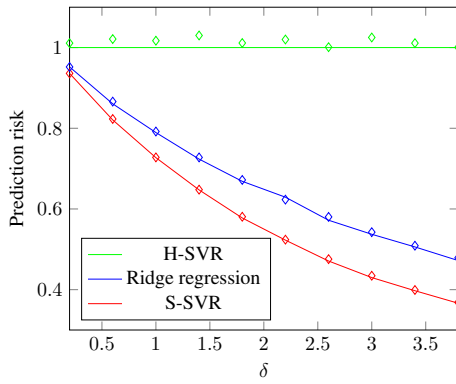
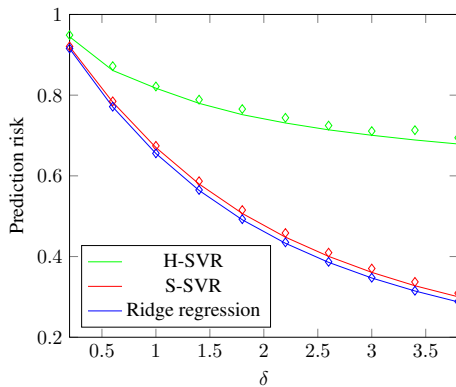

(a) $d = 3$

(b) $d = 10$

Figure 8. Prediction risk of S-SVR and ridge regression vs. δ when optimal regularization is used. Comparison of S-SVR and ridge regression with elliptic noise. The continuous line curves correspond to the asymptotic performance while the points denote finite-sample performance when $p = 200$ and $n = \lfloor \delta p \rfloor$.

Belkin, M., Hsu, D., and Xu, J. Two models of double descent for weak features. *ArXiv*, abs/1903.07571, 2019b.

Belkin, M., Rakhlin, A., and Tsybakov, A. B. Does data interpolation contradict statistical optimality? In Chaudhuri, K. and Sugiyama, M. (eds.), *Proceedings of Machine Learning Research*, volume 89 of *Proceedings of Machine Learning Research*, pp. 1611–1619. PMLR, 16–18 Apr 2019c.

Bös, S. and Oppel, M. Dynamics of training. *Advances in Neural Information Processing Systems*, pp. 141–147, 1997.

Boyd, S. and Vandenberghe, L. *Convex Optimization*. Cambridge University Press, 2003.

Candes, E. J. and Sur, P. The phase transition for the existence of the maximum likelihood estimate in high-dimensional logistic regression. *arXiv e-prints*, art. arXiv:1804.09753, April 2018.

Cherkassky, V. and Ma, Y. Selection of meta-parameters for support vector regression. In *Artificial Neural Networks — ICANN 2002*, pp. 687–693, Berlin, Heidelberg, 2002. Springer Berlin Heidelberg.

Cherkassky, V. and Ma, Y. Practical selection of svm parameters and noise estimation for svm regression. *Neural networks*, 17(1):113–126, 2004.

Deng, Z., Kammoun, A., and Thrampoulidis, C. A model of double descent for high-dimensional binary linear classification. *submitted to Information and Inference Journal of the IMA*, 2020.

Geiger, M., Jacot, A., Spigler, S., Gabriel, F., Sagun, L., d’Ascoli, S., Biroli, G., Hongler, C., and Wyart, M. Scaling description of generalization with number of parameters in deep learning. *ArXiv*, abs/1901.01608, 2019.

Gordon, Y. Some inequalities for gaussian processes and applications. *Israel Journal of Mathematics*, 50(4):265–289, dec 1985.

Gordon, Y. *On Milman’s inequality and random subspaces which escape through a mesh in $\mathbb{R}^n$* . Springer, 1988.

Guajardo, J., Weber, R., and Miranda, J. A forecasting methodology using support vector regression and dynamic feature selection. *Journal of Information & Knowledge Management*, 5(04):329–335, 2006.

Hable, R. and Christmann, A. On qualitative robustness of support vector machines. *Journal of Multivariate Analysis*, 102(6):993–1007, 2011.

Hamdi, T., Ali, J. B., Di Costanzo, V., Fnaiech, F., Moreau, E., and Ginoux, J.-M. Accurate prediction of continuous blood glucose based on support vector regression and differential evolution algorithm. *Biocybernetics and Biomedical Engineering*, 38(2):362–372, 2018.

Hastie, T., Montanari, A., Rosset, S., and Tibshirani, R. J. Surprises in high-dimensional ridgeless least squares interpolation. *arXiv preprint arXiv:1903.08560*, 2019.

Ju, X., Cheng, M., Xia, Y., Quo, F., and Tian, Y. Support vector regression and time series analysis for the forecasting of bayannur’s total water requirement. *Procedia Computer Science*, 31:523–531, 2014.

Kammoun, A. and Alouini, M.-S. On the precise error analysis of support vector machines. *arXiv preprint arXiv:2003.12972*, 2020a.

Kammoun, A. and Alouini, M.-S. On the precise error analysis of support vector machines. *Submitted to Open journal of signal processing*, 2020b.

- Kini, G. R. and Thrampoulidis, C. Analytic study of double descent in binary classification: The impact of loss. In *2020 IEEE International Symposium on Information Theory (ISIT)*, pp. 2527–2532, 2020. doi: 10.1109/ISIT44484.2020.9174344.
- Kobak, D., Lomond, J., and Sanchez, B. The optimal ridge penalty for real-world high-dimensional data can be zero or negative due to the implicit ridge regularization. *J. Mach. Learn. Res.*, 21:169:1–169:16, 2020.
- Lolas, P. Regularization in high-dimensional regression and classification via random matrix theory. *arXiv: Statistics Theory*, 2020.
- Mitra, P. P. Understanding overfitting peaks in generalization error: Analytical risk curves for l_2 and l_1 penalized interpolation. *arXiv:1906.03667*, 2019.
- Montanari, A., Ruan, F., Sohn, Y., and Yan, J. The generalization error of max-margin linear classifiers: High-dimensional asymptotics in the overparametrized regime, 2020.
- Muthukumar, V., Vodrahalli, K., Subramanian, V., and Sahai, A. Harmless interpolation of noisy data in regression. *IEEE Journal on Selected Areas in Information Theory*, 1(1):67–83, 2020. doi: 10.1109/JSAIT.2020.2984716.
- Nakkiran, P., Kaplun, G., Bansal, Y., Yang, T., Barak, B., and Sutskever, I. Deep double descent: Where bigger models and more data hurt. In *International Conference on Learning Representations*, 2020a. URL <https://openreview.net/forum?id=Blg5sA4twr>.
- Nakkiran, P., Venkat, P., Kakade, S. M., and Ma, T. Optimal regularization can mitigate double descent. *CoRR*, abs/2003.01897, 2020b. URL <https://arxiv.org/abs/2003.01897>.
- Qu, H. and Zhang, Y. A new kernel of support vector regression for forecasting high-frequency stock returns. *Mathematical Problems in Engineering*, 2016.
- Sion, M. On general minimax theorems. *Pacific Journal of Mathematics*, 8:171–176, 1958.
- Spigler, S., Geiger, M., d’Ascoli, S., Sagun, L., Biroli, G., and Wyart, M. A jamming transition from under- to over-parametrization affects generalization in deep learning. *Journal of Physics A: Mathematical and Theoretical*, 52, 10 2019. doi: 10.1088/1751-8121/ab4c8b.
- Thrampoulidis, C., Abbasi, E., and Hassibi, B. Precise Error Analysis of Regularized M-Estimators in High Dimensions. *IEEE Transactions on Information Theory*, 64(8), August 2018.
- Xu, H., Caramanis, C., and Mannor, S. Robustness and regularization of support vector machines. *Journal of machine learning research*, 10(7), 2009.
- Yang, H., Huang, K., King, I., and Lyu, M. R. Localized support vector regression for time series prediction. *Neurocomputing*, 72(10-12):2659–2669, 2009.

6. CGMT framework

The proofs of our results are based on the extension of the CGMT framework, recently proposed in (Deng et al., 2020), that allows for handling optimization problems in which the optimization sets may not be compact or are not necessarily feasible. For the reader convenience, we provide a review of the main results (Deng et al., 2020) that will be extensively used in our proofs.

In order to summarize the essential ideas, we consider the following two Gaussian processes:

$$X_{\mathbf{w}, \mathbf{u}} := \mathbf{u}^T \mathbf{G} \mathbf{w} + \psi(\mathbf{w}, \mathbf{u}) \quad (16)$$

$$Y_{\mathbf{w}, \mathbf{u}} := \|\mathbf{w}\|_2 \mathbf{g}^T \mathbf{u} + \|\mathbf{u}\|_2 \mathbf{h}^T \mathbf{w} + \psi(\mathbf{w}, \mathbf{u}) \quad (17)$$

where $\mathbf{G} \in \mathbb{R}^{n \times d}$, $\mathbf{g} \in \mathbb{R}^n$, $\mathbf{h} \in \mathbb{R}^d$ have standard normal i.i.d. entries and $\psi : \mathbb{R}^d \times \mathbb{R}^n \rightarrow \mathbb{R}$ is a continuous function. Let $\mathcal{S}_{\mathbf{w}} \subset \mathbb{R}^d$ and $\mathcal{S}_{\mathbf{u}} \subset \mathbb{R}^n$ be two convex sets but not necessarily compact. We consider the following random min-max optimization problems which refer to as the primary optimization (PO) problem and the auxiliary optimization problem (AO):

$$(\text{PO}) \quad \Phi(\mathbf{G}) = \min_{\mathbf{w} \in \mathcal{S}_{\mathbf{w}}} \max_{\mathbf{u} \in \mathcal{S}_{\mathbf{u}}} X_{\mathbf{w}, \mathbf{u}} \quad (18)$$

$$(\text{AO}) \quad \phi(\mathbf{g}, \mathbf{h}) = \min_{\mathbf{w} \in \mathcal{S}_{\mathbf{w}}} \max_{\mathbf{u} \in \mathcal{S}_{\mathbf{u}}} Y_{\mathbf{w}, \mathbf{u}} \quad (19)$$

The direct analysis of the (PO) is in general difficult. The power of the CGMT is that it transfers the asymptotic behavior of the PO to that of the AO. Particularly, if $\mathcal{S}_{\mathbf{w}}$ and $\mathcal{S}_{\mathbf{u}}$ are compact sets, the use of Gordon's inequality (Gordon, 1988; 1985) implies that:

$$\mathbb{P}(\Phi(\mathbf{G}) < c) \leq 2\mathbb{P}(\phi(\mathbf{g}, \mathbf{h}) < c) \quad (20)$$

The only conditions required by the above inequality is that the sets are compact and ψ is continuous. Particularly, Gordon's inequality holds even if the optimization sets are not convex and ψ is not convex-concave. As a by product of (20), it follows that if c is a high-probability lower bound of the AO then it is also a lower bound of the PO. If the set $\mathcal{S}_{\mathbf{w}}$ and $\mathcal{S}_{\mathbf{u}}$ are additionally convex and ψ is convex-concave then, for any $\nu \in \mathbb{R}$ and $t > 0$, it holds

$$\mathbb{P}(|\Phi(\mathbf{G}) - \nu| > t) \leq 2\mathbb{P}(|\phi(\mathbf{g}, \mathbf{h}) - \nu| > t)$$

In other words, if the AO cost concentrates around ν^* then the optimal cost of the PO concentrates also around the same value ν^* . Moreover, it has been shown in (Thrapoulidis et al., 2018) that under strict convexity conditions of the asymptotic AO, the concentration of the optimal solution of the AO translates into the concentration of the optimal solution of the PO around the same value. More precisely, denote by $\mathbf{w}_\phi(\mathbf{g}, \mathbf{h})$ the optimal solution of the AO. If one can prove that the minimizers of the AO satisfy $\|\mathbf{w}_\phi(\mathbf{g}, \mathbf{h})\| \xrightarrow{a.s.} \alpha_*$ where α_* is the unique solution of a certain limiting optimization problem whose cost is asymptotically equivalent to the AO. Then, the solution of the PO, denoted by $\mathbf{w}_\Phi(\mathbf{G})$ also satisfies $\|\mathbf{w}_\Phi(\mathbf{G})\| \xrightarrow{a.s.} \alpha_*$.

In (Deng et al., 2020), the authors developed a principled machinery that extends the results of the CGMT to problems in which the optimization sets are not compact and may not be necessarily feasible. Before reviewing the results of (Deng et al., 2020), we shall introduce some important notations. For fixed R and Γ , we consider the following “ (R, Γ) -bounded” version of the PO:

$$\Phi_{R, \Gamma}(\mathbf{G}) = \min_{\substack{\mathbf{w} \in \mathcal{S}_{\mathbf{w}} \\ \|\mathbf{w}\|_2 \leq R}} \max_{\substack{\mathbf{u} \in \mathcal{S}_{\mathbf{u}} \\ \|\mathbf{u}\|_2 \leq \Gamma}} X_{\mathbf{w}, \mathbf{u}}$$

with which we associate the following (R, Γ) bounded version of the AO:

$$\phi_{R, \Gamma}(\mathbf{g}, \mathbf{h}) = \min_{\substack{\mathbf{w} \in \mathcal{S}_{\mathbf{w}} \\ \|\mathbf{w}\|_2 \leq R}} \max_{\substack{\mathbf{u} \in \mathcal{S}_{\mathbf{u}} \\ \|\mathbf{u}\|_2 \leq \Gamma}} Y_{\mathbf{w}, \mathbf{u}} \quad (21)$$

Since $\Gamma \mapsto \max_{\substack{\mathbf{u} \in \mathcal{S}_{\mathbf{u}} \\ \|\mathbf{u}\|_2 \leq \Gamma}} X_{\mathbf{w}, \mathbf{u}}$ is concave in Γ , we can establish using Sion's min-max theorem (Sion, 1958) that

$$\Phi(\mathbf{G}) = \inf_{R \geq 0} \sup_{\Gamma \geq 0} \Phi_{R, \Gamma}$$

Similarly, we define $\phi_R(\mathbf{g}, \mathbf{h})$ as follows:

$$\phi_R(g, h) := \sup_{\Gamma \geq 0} \phi_{R,\Gamma}(\mathbf{g}, \mathbf{h}) \quad (22)$$

However, since $Y_{\mathbf{w},\mathbf{u}}$ is not convex-concave in $(\mathbf{w}, \mathbf{u})$, we only have:

$$\sup_{\Gamma \geq 0} \phi_{R,\Gamma}(\mathbf{g}, \mathbf{h}) \leq \min_{\substack{\mathbf{w} \in \mathcal{S}_{\mathbf{w}} \\ \|\mathbf{w}\|_2 \leq R}} \max_{\mathbf{u} \in \mathcal{S}_{\mathbf{u}}} Y_{\mathbf{w},\mathbf{u}}$$

The extension of the CGMT allows us to connect properties of the unbounded PO to those of the sequence of AO problems $\phi_R(\mathbf{g}, \mathbf{h})$. Theorem 4 illustrates how the study of the cost of the sequence of AO problems can help determine the feasibility of the original PO problem, whereas Theorem 5 shows how the asymptotic behavior of the optimal cost of the original problem boils down to analyzing that of the sequence of AO problems.

Theorem 4 (Feasibility (Deng et al., 2020)). *Recall the definitions of $\Phi(\mathbf{G})$, $\phi_{R,\Gamma}(\mathbf{g}, \mathbf{h})$ and $\phi_R(\mathbf{g}, \mathbf{h})$ in (19), (21) and (22), respectively. Assume that $(\mathbf{w}, \mathbf{u}) \mapsto X_{\mathbf{w},\mathbf{u}}$ in (16) is convex-concave and that the constraint sets $\mathcal{S}_{\mathbf{w}}, \mathcal{S}_{\mathbf{u}}$ are convex (but not necessarily bounded). The following two statements hold true.*

- (i) *Assume that for any fixed $R, \Gamma \geq 0$ there exists a positive constant C (independent of R and Γ) and a continuous increasing function $f : \mathbb{R}_+ \rightarrow \mathbb{R}$ tending to infinity such that, for n sufficiently large (independent of R and Γ):*

$$\phi_{R,\Gamma}(\mathbf{g}, \mathbf{h}) \geq C f(\Gamma). \quad (23)$$

Then, with probability 1, for n sufficiently large, $\Phi(\mathbf{G}) = \infty$.

- (ii) *Assume that there exists $k_0 \in \mathbb{N}$ and a positive constant C such that:*

$$\mathbb{P}[\{\phi_{k_0}(\mathbf{g}, \mathbf{h}) \geq C\}, \text{ i.o.}] = 0. \quad (24)$$

Then, $\mathbb{P}[\Phi(\mathbf{G}) = \infty, \text{ i.o.}] = 0$.

Theorem 5. (Deng et al., 2020) *Assume the same notation as in Theorem 4. Assume that there exists $\bar{\phi}$ and $k_0 \in \mathbb{N}$ such that the following statements hold true:*

$$\text{For any } \epsilon > 0 : \quad \mathbb{P}\left[\bigcup_{k=k_0}^{\infty} \left\{\phi_k(\mathbf{g}, \mathbf{h}) \leq \bar{\phi} - \epsilon\right\}, \text{ i.o.}\right] = 0, \quad (25a)$$

$$\text{For any } \epsilon > 0 : \quad \mathbb{P}\left[\left\{\phi_{k_0}(\mathbf{g}, \mathbf{h}) \geq \bar{\phi} + \epsilon\right\}, \text{ i.o.}\right] = 0, \quad (25b)$$

$$(25c)$$

Then,

$$\Phi(\mathbf{G}) \xrightarrow{a.s.} \bar{\phi}. \quad (26)$$

Further, let $\mathcal{S}$ be an open subset of $\mathcal{S}_{\mathbf{w}}$ and $\mathcal{S}^c = \mathcal{S}_{\mathbf{w}} \setminus \mathcal{S}$. Denote by $\tilde{\phi}_{R,\Gamma}(\mathbf{g}, \mathbf{h})$ the optimal cost of (21) when the minimization over $\mathbf{w}$ is now further constrained over $\mathbf{w} \in \mathcal{S}^c$. Define $\tilde{\phi}_R(\mathbf{g}, \mathbf{h})$ in a similar way to (22). Assume the following statement hold true,

$$\text{There exists } \zeta > 0 : \quad \mathbb{P}\left[\bigcup_{k=k_0}^{\infty} \left\{\tilde{\phi}_k(\mathbf{g}, \mathbf{h}) \leq \bar{\phi} + \zeta\right\}, \text{ i.o.}\right] = 0. \quad (27)$$

Then, letting $\mathbf{w}_{\Phi}$ denote a minimizer of the PO in (18), it also holds that

$$\mathbb{P}[\mathbf{w}_{\Phi} \in \mathcal{S}, \text{ for sufficiently large } n] = 1. \quad (28)$$

7. Proof of Theorem 1 and Theorem 2

In this section, we analyze the statistical behavior of H-SVR.

7.1. Preliminaries

Identification of the PO: The Lagrangian function associated with H-SVR is given by:

$$\mathcal{L}(\mathbf{w}, \boldsymbol{\lambda}, \boldsymbol{\alpha}) := \frac{1}{2} \|\mathbf{w}\|^2 + \sum_{i=1}^n \lambda_i (\boldsymbol{\beta}_*^T \mathbf{x}_i + \sigma n_i - \mathbf{w}^T \mathbf{x}_i - \epsilon) + \sum_{i=1}^n \alpha_i (\mathbf{w}^T \mathbf{x}_i - \boldsymbol{\beta}_*^T \mathbf{x}_i - \sigma n_i - \epsilon)$$

where $\boldsymbol{\lambda} = [\lambda_1, \dots, \lambda_n]$ and $\boldsymbol{\alpha} = [\alpha_1, \dots, \alpha_n]$ are the Lagrange multipliers. From the first order optimality conditions, we have:

$$\frac{\partial \mathcal{L}}{\partial \mathbf{w}} = 0 \implies \mathbf{w} = \sum_{i=1}^n (\lambda_i - \alpha_i) \mathbf{x}_i \quad (29)$$

Leveraging the above relations, the optimization problem becomes equivalent to solving:

$$\max_{\boldsymbol{\lambda} \geq 0, \boldsymbol{\alpha} \geq 0} -\frac{1}{2} (\boldsymbol{\lambda} - \boldsymbol{\alpha})^T \mathbf{X}^T \mathbf{X} (\boldsymbol{\lambda} - \boldsymbol{\alpha}) + \sigma \mathbf{n}^T (\boldsymbol{\lambda} - \boldsymbol{\alpha}) + \boldsymbol{\beta}_*^T \mathbf{X} (\boldsymbol{\lambda} - \boldsymbol{\alpha}) - \epsilon \mathbf{1}^T (\boldsymbol{\lambda} + \boldsymbol{\alpha})$$

where $\mathbf{X} = [\mathbf{x}_1, \dots, \mathbf{x}_n]$. Performing the change of variables $\mathbf{u} = \sqrt{p}(\boldsymbol{\lambda} - \boldsymbol{\alpha})$ and $\mathbf{v} = \sqrt{p}(\boldsymbol{\lambda} + \boldsymbol{\alpha})$, the above problem simplifies to:

$$\max_{\substack{\mathbf{u}, \mathbf{v} \\ |\mathbf{u}| \leq \mathbf{v} \\ \mathbf{v} \geq 0}} -\frac{1}{2p} \mathbf{u}^T \mathbf{X}^T \mathbf{X} \mathbf{u} + \frac{1}{\sqrt{p}} \sigma \mathbf{n}^T \mathbf{u} + \frac{1}{\sqrt{p}} \boldsymbol{\beta}_*^T \mathbf{X} \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T \mathbf{v} \quad (30)$$

From (29), the solution of the H-SVR is related to an optimal solution $\hat{\mathbf{u}}$ of (30) as:

$$\hat{\mathbf{w}}_H = \frac{1}{\sqrt{p}} \mathbf{X} \hat{\mathbf{u}} \quad (31)$$

Obviously, for a given $\mathbf{u}$ the optimum over $\mathbf{v}$ of the objective cost in (30) is given by $\mathbf{v} = |\mathbf{u}|$. Replacing $\mathbf{v}$ by its optimal value, we obtain

$$\max_{\mathbf{u}} -\frac{1}{2p} \mathbf{u}^T \mathbf{X}^T \mathbf{X} \mathbf{u} + \frac{1}{\sqrt{p}} \sigma \mathbf{n}^T \mathbf{u} + \frac{1}{\sqrt{p}} \boldsymbol{\beta}_*^T \mathbf{X} \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T |\mathbf{u}| \quad (32)$$

To write the above problem in the form required by the CGMT, we shall use the following relation:

$$-\frac{1}{2p} \mathbf{u}^T \mathbf{X}^T \mathbf{X} \mathbf{u} + \frac{1}{\sqrt{p}} \boldsymbol{\beta}_*^T \mathbf{X} \mathbf{u} = \min_{\tilde{\mathbf{w}}} \frac{1}{\sqrt{p}} \tilde{\mathbf{w}}^T \mathbf{X} \mathbf{u} + \frac{1}{2} \|\tilde{\mathbf{w}}\|^2 - \boldsymbol{\beta}_*^T \tilde{\mathbf{w}} + \frac{1}{2} \|\boldsymbol{\beta}_*\|^2, \quad (33)$$

where in the right-hand side term, the optimal $\hat{\tilde{\mathbf{w}}}_H$ can be easily verified to be unique and given by:

$$\hat{\tilde{\mathbf{w}}}_H = -\frac{1}{\sqrt{p}} \mathbf{X} \hat{\mathbf{u}} + \boldsymbol{\beta}_* \quad (34)$$

Plugging (33) into (30), we identify the following problem which takes the form of a primary optimization problem as required by the CGTM framework:

$$\Phi^{(n)} = \max_{\mathbf{u}} \min_{\tilde{\mathbf{w}}} \frac{1}{\sqrt{p}} \mathbf{u}^T \mathbf{X} \tilde{\mathbf{w}} + \frac{1}{2} \|\tilde{\mathbf{w}}\|^2 - \boldsymbol{\beta}_*^T \tilde{\mathbf{w}} + \sigma \frac{1}{\sqrt{p}} \mathbf{n}^T \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T |\mathbf{u}| + \frac{1}{2} \|\boldsymbol{\beta}_*\|^2 \quad (35)$$

From (31) and (34) and due to the uniqueness of $\hat{\tilde{\mathbf{w}}}_H$ in (33), the H-SVR solution satisfies:

$$\hat{\tilde{\mathbf{w}}}_H = -\hat{\mathbf{w}}_H + \boldsymbol{\beta}_* \quad (36)$$

Let $\mathbf{P}_\perp = \mathbf{I}_p - \frac{\boldsymbol{\beta}_* \boldsymbol{\beta}_*^T}{\boldsymbol{\beta}_*^T \boldsymbol{\beta}_*}$. Then, we can decompose $\mathbf{u}^T \mathbf{X}^T \tilde{\mathbf{w}}$ as:

$$\mathbf{u}^T \mathbf{X}^T \tilde{\mathbf{w}} = \mathbf{u}^T \mathbf{X}^T \mathbf{P}_\perp \tilde{\mathbf{w}} + \frac{1}{\|\boldsymbol{\beta}_*\|^2} \mathbf{u}^T \mathbf{X}^T \boldsymbol{\beta}_* \boldsymbol{\beta}_*^T \tilde{\mathbf{w}}$$

Plugging this decomposition into (35), we obtain:

$$\Phi^{(n)} = \max_{\mathbf{u}} \min_{\tilde{\mathbf{w}}} \frac{1}{\sqrt{p}} \mathbf{u}^T \mathbf{X}^T \mathbf{P}_{\perp} \tilde{\mathbf{w}} + \frac{1}{\sqrt{p}} \frac{1}{\|\beta_{\star}\|^2} \mathbf{u}^T \mathbf{X}^T \beta_{\star} \beta_{\star}^T \tilde{\mathbf{w}} + \frac{1}{2} \|\beta_{\star} - \tilde{\mathbf{w}}\|^2 + \frac{1}{\sqrt{p}} \sigma \mathbf{n}^T \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T |\mathbf{u}|$$

Let $\mathbf{z} = \frac{1}{\|\beta_{\star}\|} \mathbf{X}^T \beta_{\star}$. Clearly, $\mathbf{z}$ is a standard Gaussian random vector and is independent of $\mathbf{X}^T \mathbf{P}_{\perp}$. Replacing $\frac{1}{\|\beta_{\star}\|} \mathbf{X}^T \beta_{\star}$ by $\mathbf{z}$, $\Phi^{(n)}$ writes as:

$$\Phi^{(n)} = \min_{\tilde{\mathbf{w}}} \max_{\mathbf{u}} \frac{1}{\sqrt{p}} \mathbf{u}^T \mathbf{X}^T \mathbf{P}_{\perp} \tilde{\mathbf{w}} + \frac{1}{\sqrt{p}} \frac{1}{\|\beta_{\star}\|} \mathbf{u}^T \mathbf{z} \beta_{\star}^T \tilde{\mathbf{w}} + \frac{1}{2} \|\beta_{\star} - \tilde{\mathbf{w}}\|^2 + \frac{1}{\sqrt{p}} \sigma \mathbf{n}^T \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T |\mathbf{u}|$$

Identification of the AO sequences: Following the notations of section 6 in the Appendix, we associate with the PO a sequence of auxiliary problems indexed by r and θ , where each one is associated with a “double bounded” version of the PO. Particularly, for fixed r and θ , we define every element of the AO sequence as:

$$\phi_{r,\theta}^{(n)} := \min_{\substack{\tilde{\mathbf{w}} \\ \|\tilde{\mathbf{w}}\| \leq r}} \max_{\substack{\mathbf{u} \\ \|\mathbf{u}\| \leq \theta}} \frac{1}{\sqrt{p}} \|\mathbf{P}_{\perp} \tilde{\mathbf{w}}\|_2 \mathbf{g}^T \mathbf{u} + \frac{1}{\sqrt{p}} \|\mathbf{u}\| \mathbf{h}^T \mathbf{P}_{\perp} \tilde{\mathbf{w}} + \frac{1}{\sqrt{p}} \frac{1}{\|\beta_{\star}\|} \mathbf{u}^T \mathbf{z} \beta_{\star}^T \tilde{\mathbf{w}} + \frac{1}{2} \|\beta_{\star} - \tilde{\mathbf{w}}\|^2 + \frac{1}{\sqrt{p}} \sigma \mathbf{n}^T \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T |\mathbf{u}|$$

Simplification of the AO problems: Having identified the AO sequence, we proceed now with some elementary manipulations to reduce each AO problem to a scalar optimization problem. For that and considering the fact that $\tilde{\mathbf{w}}$ intervenes in the objective through its norm, its projection along $\beta_{\star}$ and the space orthogonal to it, it is sensible to decompose it as:

$$\tilde{\mathbf{w}} = \gamma_1 \frac{\mathbf{P}_{\perp} \tilde{\mathbf{w}}}{\|\mathbf{P}_{\perp} \tilde{\mathbf{w}}\|} + \gamma_2 \frac{\beta_{\star}}{\|\beta_{\star}\|}$$

Using this decomposition, the AO problem simplifies to:

$$\phi_{r,\theta}^{(n)} = \min_{\gamma_1^2 + \gamma_2^2 \leq r^2} \max_{0 \leq m \leq \theta} \max_{\|\mathbf{u}\|_2 = m} \frac{1}{2} \gamma_1^2 + \frac{1}{2} \gamma_2^2 + \gamma_1 \frac{1}{\sqrt{p}} \mathbf{g}^T \mathbf{u} - m \gamma_1 \frac{1}{\sqrt{p}} \|\mathbf{P}_{\perp} \mathbf{h}\| + \gamma_2 \frac{1}{\sqrt{p}} \mathbf{u}^T \mathbf{z} + \frac{1}{2} \|\beta_{\star}\|^2 - \gamma_2 \|\beta_{\star}\| \quad (37)$$

$$+ \frac{1}{\sqrt{p}} \sigma \mathbf{n}^T \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T |\mathbf{u}| \quad (38)$$

Hence,

$$\begin{aligned} \phi_{r,\theta}^{(n)} = \min_{\gamma_1^2 + \gamma_2^2 \leq r^2} \max_{0 \leq m \leq \theta} \max_{\|\mathbf{u}\|_2 = m} \frac{1}{2} \gamma_1^2 + \frac{1}{2} \gamma_2^2 + \frac{1}{\sqrt{p}} \mathbf{u}^T (\gamma_1 \mathbf{g} + \gamma_2 \mathbf{z} + \sigma \mathbf{n}) - \frac{\epsilon}{\sqrt{p}} \|\mathbf{u}\|_1 - m \gamma_1 \frac{1}{\sqrt{p}} \|\mathbf{P}_{\perp} \mathbf{h}\|_2 \\ + \frac{1}{2} \|\beta_{\star}\|_2^2 - \gamma_2 \|\beta_{\star}\| \end{aligned} \quad (39)$$

Using Lemma 1, we obtain:

$$\begin{aligned} \phi_{r,\theta}^{(n)} = \min_{\gamma_1^2 + \gamma_2^2 \leq r^2} \max_{0 \leq m \leq \theta} \frac{1}{2} \gamma_1^2 + \frac{1}{2} \gamma_2^2 + m \left(\sqrt{\frac{1}{p} \sum_{i=1}^n (|\gamma_1 g_i + \gamma_2 z_i + \sigma n_i| - \epsilon)_+^2} - \gamma_1 \frac{1}{\sqrt{p}} \|\mathbf{P}_{\perp} \mathbf{h}\|_2 \right) \\ + \frac{1}{2} \|\beta_{\star}\|_2^2 - \gamma_2 \|\beta_{\star}\|_2 \\ = \min_{\gamma_1^2 + \gamma_2^2 \leq r^2} \frac{1}{2} \gamma_1^2 + \frac{1}{2} \gamma_2^2 + \theta \left(\left(\sqrt{\frac{1}{p} \sum_{i=1}^n (|\gamma_1 g_i + \gamma_2 z_i + \sigma n_i| - \epsilon)_+^2} - \gamma_1 \frac{1}{\sqrt{p}} \|\mathbf{P}_{\perp} \mathbf{h}\|_2 \right) \right)_+ \\ + \frac{1}{2} \|\beta_{\star}\|_2^2 - \gamma_2 \|\beta_{\star}\|_2 \end{aligned} \quad (40)$$

At this point, it is worth pointing out that the new formulation of the AO sequence problems is reduced to optimization problems that involves only few scalar variables. It relies on a deterministic analysis that does not involve any asymptotic approximation.

Asymptotic behavior of the AOs Let us consider the following sequence of functions:

$$D_n(\gamma_1, \gamma_2) := \left\| \left(\left| \frac{\gamma_1}{\sqrt{p}} \mathbf{g} + \frac{\gamma_2}{\sqrt{p}} \mathbf{z} + \sigma \mathbf{n} \right| - \epsilon \right)_+ \right\|_2 - \frac{\gamma_1}{\sqrt{p}} \|\mathbf{P}_\perp \mathbf{h}\|_2$$

defined for γ_1 and γ_2 such that $\gamma_1^2 + \gamma_2^2 \leq r^2$. It is easy to note that $(\gamma_1, \gamma_2) \mapsto D_n(\gamma_1, \gamma_2)$ is jointly convex in its arguments (γ_1, γ_2) and by (Thrapoulidis et al., 2018, Lemma 10) converges point-wise to

$$(\gamma_1, \gamma_2) \mapsto \bar{D}(\gamma_1, \gamma_2)$$

with

$$\bar{D}(\gamma_1, \gamma_2) := \sqrt{\delta} \sqrt{\mathbb{E} \left(\left| \sqrt{\gamma_1^2 + \gamma_2^2} G + \sigma N \right| - \epsilon \right)_+^2} - \gamma_1$$

where $G \sim \mathcal{N}(0, 1)$. Since convergence of convex functions is uniform over compacts, then letting $\mathcal{C}$ be an arbitrary compact in $\mathbb{R}^2$, and any $\eta > 0$, for n sufficiently large, it holds that:

$$\bar{D}(\gamma_1, \gamma_2) - \eta \leq D_n(\gamma_1, \gamma_2) \leq \bar{D}(\gamma_1, \gamma_2) + \eta, \quad \forall (\gamma_1, \gamma_2) \in \mathcal{C} \quad (41)$$

7.2. Proof of Theorem 1

Letting $\delta_\star = \frac{1}{\inf_{t \in \mathbb{R}} \mathbb{E} (|G + t\sigma N| - t\epsilon)_+^2}$, we will prove the following two statements:

$$\delta > \delta_\star \Rightarrow \mathbb{P} [\text{The hard-margin SVR is feasible for sufficiently large } n] = 0 \quad (42)$$

$$\delta < \delta_\star \Rightarrow \mathbb{P} [\text{The hard-margin SVR is feasible for sufficiently large } n] = 1. \quad (43)$$

In the sequel, we will sequentially establish (42) and (43).

Proof of (42). To prove the desired result, we will apply Theorem 4. More precisely, in view of (23) in Theorem 4, it suffices to prove that under the condition $\delta > \delta_\star$, for any fixed r and θ , there exists constant $C > 0$ that is independent of r and θ such that for sufficiently large n (taken independently of r and θ):

$$\phi_{r,\theta}^{(n)} \geq C\theta \quad (44)$$

To prove (44), we start by noting that:

$$\phi_{r,\theta}^{(n)} = \min_{\gamma_1^2 + \gamma_2^2 \leq r^2} \frac{1}{2} (\gamma_2 - \|\beta_\star\|)^2 + \frac{1}{2} \gamma_1^2 + \theta \left(D_n(\gamma_1, \gamma_2) \right)_+ \quad (45)$$

Hence,

$$\phi_{r,\theta}^{(n)} \geq \theta \left(\min_{\gamma_1, \gamma_2} D_n(\gamma_1, \gamma_2) \right)_+ \quad (46)$$

Function $(\gamma_1, \gamma_2) \mapsto D_n(\gamma_1, \gamma_2)$ is jointly convex in its arguments (γ_1, γ_2) and converges pointwise to $(\gamma_1, \gamma_2) \mapsto \bar{D}(\gamma_1, \gamma_2)$. Using Lemma 3, for γ_1 fixed, $\gamma_2 \mapsto D_n(\gamma_1, \gamma_2)$ is convex. It converges pointwise to $\gamma_2 \mapsto \bar{D}(\gamma_1, \gamma_2)$. As $\lim_{\gamma_2 \rightarrow \pm\infty} \bar{D}(\gamma_1, \gamma_2) = \infty$, we may use (Thrapoulidis et al., 2018, Lemma 10) to obtain,

$$\min_{\gamma_2 \in \mathbb{R}} D_n(\gamma_1, \gamma_2) \xrightarrow{a.s.} \min_{\gamma_2 \in \mathbb{R}} \bar{D}(\gamma_1, \gamma_2)$$

Moreover, it is easy to see that $\gamma_2 \mapsto \bar{D}(\gamma_1, \gamma_2)$ is an even function. As it is convex, from Lemma 2, the minimum is taken at $\gamma_2 = 0$ and hence $\min_{\gamma_2} \bar{D}(\gamma_1, \gamma_2) = \bar{D}(\gamma_1, 0)$. Similarly, $\gamma_1 \mapsto \min_{\gamma_2} D_n(\gamma_1, \gamma_2)$ is convex and converges pointwise to $\gamma_1 \mapsto \bar{D}(\gamma_1, 0)$. Moreover, for fixed γ_1 , it is easy to see that

$$\bar{D}(\gamma_1, 0) = |\gamma_1| \left(\sqrt{\delta} \sqrt{\mathbb{E} \left(\left| G + \frac{\sigma}{\gamma_1} N \right| - \frac{1}{|\gamma_1|} \epsilon \right)_+^2} - \frac{\gamma_1}{|\gamma_1|} \right) \quad (47)$$

$$\geq |\gamma_1| \left(\sqrt{\delta} \sqrt{\inf_{t \in \mathbb{R}} \mathbb{E} (|G + t\sigma N| - t\epsilon)_+^2} - 1 \right) \quad (48)$$

Since by assumption $\delta > \delta_*$,

$$\inf_{t \in \mathbb{R}} \sqrt{\delta} \mathbb{E} (|G + t\sigma N| - t\epsilon)_+^2 > 1$$

Hence, $\lim_{\gamma_1 \rightarrow \pm\infty} \min_{\gamma_2} \overline{D}(\gamma_1, \gamma_2) = \lim_{\gamma_1 \rightarrow \infty} \overline{D}(\gamma_1, 0) = \infty$. Using again (Thrampoulidis et al., 2018, Lemma 10), we obtain:

$$\min_{\gamma_1, \gamma_2} D_n(\gamma_1, \gamma_2) \xrightarrow{a.s.} \min_{\gamma_1} \overline{D}(\gamma_1, 0)$$

It follows from (46) that for any $\eta > 0$, the following holds true once n is taken sufficiently large independently of r and θ ,

$$\phi_{r,\theta}^{(n)} \geq \theta \left(\min_{\gamma_1} \overline{D}(\gamma_1, 0) - \eta \right)_+$$

Assume that

$$\min_{\gamma_1} \overline{D}(\gamma_1, 0) > 2C. \quad (49)$$

where C is some strictly positive constant. Then we can choose η sufficiently small, for instance smaller than C , to obtain (44). To conclude, it suffices thus to establish (49). To this end, first notice that for all γ_1 such that $|\gamma_1| \leq \tilde{\eta} := \frac{1}{2} \sqrt{\delta \mathbb{E} [(\sigma|N| - \epsilon)_+^2]}$, the following holds true:

$$\min_{\substack{\gamma_1 \\ |\gamma_1| \leq \tilde{\eta}}} \overline{D}(\gamma_1, 0) \geq \min_{\substack{\gamma_1 \\ |\gamma_1| \leq \tilde{\eta}}} \sqrt{\delta \mathbb{E} (|\gamma_1 G + \sigma N| - \epsilon)_+^2} - \tilde{\eta} \quad (50)$$

$$\geq \sqrt{\delta \mathbb{E} (|\sigma N| - \epsilon)_+^2} - \tilde{\eta} \quad (51)$$

$$\geq \frac{1}{2} \sqrt{\delta \mathbb{E} (|\sigma N| - \epsilon)_+^2} = \tilde{\eta} > 0. \quad (52)$$

where the last inequality follows by Lemma 2. On the other hand,

$$\min_{\substack{\gamma_1 \\ |\gamma_1| \geq \tilde{\eta}}} \overline{D}(\gamma_1, 0) \geq \min_{\substack{\gamma_1 \\ |\gamma_1| \geq \tilde{\eta}}} |\gamma_1| \left(\sqrt{\delta \mathbb{E} \left(\left| G + \frac{\sigma}{\gamma_1} N \right| - \frac{1}{|\gamma_1|} \epsilon \right)_+^2} - 1 \right) \quad (53)$$

$$\geq \tilde{\eta} \left(\sqrt{\delta \inf_{t \in \mathbb{R}} \mathbb{E} (|G + t\sigma N| - t\epsilon)_+^2} - 1 \right) = \tilde{\eta} \left(\sqrt{\frac{\delta}{\delta_*}} - 1 \right) > 0. \quad (54)$$

Putting (52) and (54) together, we obtain:

$$\min_{\gamma_1} \overline{D}(\gamma_1, 0) \geq \min \left(\tilde{\eta}, \tilde{\eta} \left(\sqrt{\frac{\delta}{\delta_*}} - 1 \right) \right)$$

Letting $C := \frac{1}{2} \min \left(\tilde{\eta}, \tilde{\eta} \left(\sqrt{\frac{\delta}{\delta_*}} - 1 \right) \right)$ yields (49).

Proof of (43) The aim here is to show that if $\delta < \delta_*$, then the PO and thus the H-SVR is feasible with probability 1 for sufficiently large n . For that, we will apply the second item of Theorem 4. To begin with, we shall start by showing that there exists $\kappa_0 > 0$ such that for all $\kappa \leq \kappa_0$ the following set is non-empty:

$$\mathcal{I} := \{\gamma_1 \mid \overline{D}(\gamma_1, 0) \leq -\kappa\} \neq \emptyset \quad (55)$$

and open. Let t_* be such that $t_* \in \arg \inf_{t \in \mathbb{R}} \mathbb{E} (|G + t\sigma N| - t\epsilon)_+^2$. Obviously t_* is finite. To see this, it suffices to note that:

$$\mathbb{E} (|G + t\sigma N| - t\epsilon)_+^2 \geq \mathbb{E} \left[(|G + t\sigma N| - t\epsilon)_+^2 \mathbf{1}_{\{\sigma N \geq 2\epsilon\}} \mathbf{1}_{\{G > 0\}} \right] \quad (56)$$

$$\geq t^2 \epsilon^2 \mathbf{1}_{\{\sigma N \geq 2\epsilon\}} \mathbf{1}_{\{G > 0\}} \quad (57)$$

$$\xrightarrow{t \rightarrow \pm\infty} \infty \quad (58)$$

By assumption $\sqrt{\delta} \sqrt{\mathbb{E}(|G + t\sigma N| - t\epsilon)_+^2} < 1$. Hence $\overline{D}(\frac{1}{\lfloor t_* \rfloor}, 0) < 0$. By continuity of $\gamma_1 \mapsto \overline{D}(\gamma_1, 0)$ we prove that $\mathcal{I}$ is non-empty and open. With this at hand, we let C_κ be given by:

$$C_\kappa := \sup_{\gamma_1 \in \mathcal{I}} |\gamma_1| \quad (59)$$

Recall that from the uniform convergence of $\gamma \mapsto D_n(\gamma, 0)$ to $\gamma \mapsto \overline{D}(\gamma, 0)$ over compacts, we may argue that for any $\gamma_1 \in [-C_\kappa, C_\kappa]$ and for n sufficiently large:

$$D_n(\gamma_1, 0) \leq \overline{D}(\gamma_1, 0) + \kappa$$

Hence,

$$\{\gamma_1, D_n(\gamma_1, 0) \leq 0\} \subset \{\overline{D}(\gamma_1, 0) \leq -\kappa\}$$

and as such:

$$\left\{ \{\gamma_1, D_n(\gamma_1, 0) \leq 0\} \cap [-C_\kappa, C_\kappa] \right\} \subset \left\{ \{\overline{D}(\gamma_1, 0) \leq -\kappa\} \cap [-C_\kappa, C_\kappa] \right\}$$

For $k \in \mathbb{R}$, define $\phi_k^{(n)}$:

$$\phi_k^{(n)} = \sup_{\theta \geq 0} \phi_{r, \theta}^{(n)}$$

Then, using the min-max inequality, we have:

$$\phi_k^{(n)} \leq \min_{\gamma_1, \gamma_2} \sup_{\theta \geq 0} \frac{1}{2}(\gamma_2 - \|\beta_\star\|)^2 + \frac{1}{2}\gamma_1^2 + \theta(D_n(\gamma_1, \gamma_2))_+ \quad (60)$$

$$\leq \min_{\substack{\gamma_1 \in \mathcal{I} \\ |\gamma_1| \leq C_\kappa}} \sup_{\theta \geq 0} \frac{1}{2}\|\beta_\star\|^2 + \frac{1}{2}\gamma_1^2 + \theta(D_n(\gamma_1, 0))_+ \quad (61)$$

$$\leq \frac{1}{2}(C_\kappa^2 + \|\beta\|^2) \quad (62)$$

This shows that there exists a constant C such that almost surely:

$$\phi_k^{(n)} \leq C.$$

By Theorem 4-ii), we conclude that if $\delta < \delta_\star$, the hard-margin SVR is almost surely feasible.

7.3. Proof of Theorem 2

The proof of Theorem 2 follows from applying the result of Theorem 5 in the Appendix. Let $\overline{\phi}$ be defined as:

$$\overline{\phi} = \min_{\overline{D}(\gamma_1, \gamma_2) \leq 0} \frac{1}{2}\gamma_1^2 + \frac{1}{2}(\gamma_2 - \beta)^2. \quad (63)$$

and denote by $\gamma_1^\star$ and $\gamma_2^\star$ its corresponding minimizers¹. We need to prove that the following statements hold true:

$$\text{For any } \epsilon > 0 : \mathbb{P} \left[\bigcup_{k=k_0}^\infty \left\{ \phi_k^{(n)} \leq \overline{\phi} - \epsilon \right\}, \text{ i.o.} \right] = 0. \quad (64a)$$

$$\text{For any } \epsilon > 0 : \mathbb{P} \left[\left\{ \phi_{k_0}^{(n)} \geq \overline{\phi} - \epsilon \right\}, \text{ i.o.} \right] = 0. \quad (64b)$$

where k_0 is some integer sufficiently large, and $\phi_k^{(n)}$ is defined as:

$$\phi_k^{(n)} = \sup_{\theta \geq 0} \phi_{k, \theta}^{(n)}$$

¹The existence and uniqueness of $\gamma_1^\star$ and $\gamma_2^\star$ follows from the fact that the objective is strictly convex and coercive.

We shall first simplify the expression of $\phi_k^{(n)}$. It follows from (45) that:

$$\phi_k^{(n)} = \sup_{\theta \geq 0} \min_{\substack{\gamma_1, \gamma_2 \\ \gamma_1^2 + \gamma_2^2 \leq k^2}} \frac{1}{2}(\gamma_2 - \|\beta_\star\|)^2 + \frac{1}{2}\gamma_1^2 + \theta(D_n(\gamma_1, \gamma_2))_+ \quad (65)$$

$$= \min_{\substack{\gamma_1, \gamma_2 \\ \gamma_1^2 + \gamma_2^2 \leq k^2}} \sup_{\theta \geq 0} \frac{1}{2}(\gamma_2 - \|\beta_\star\|)^2 + \frac{1}{2}\gamma_1^2 + \theta(D_n(\gamma_1, \gamma_2))_+ \quad (66)$$

$$= \min_{\substack{\gamma_1, \gamma_2 \\ \gamma_1^2 + \gamma_2^2 \leq k^2 \\ D_n(\gamma_1, \gamma_2) \leq 0}} \frac{1}{2}(\gamma_2 - \|\beta_\star\|)^2 + \frac{1}{2}\gamma_1^2 \quad (67)$$

The second equality (66) is because $(\gamma_1, \gamma_2, \theta) \mapsto \theta D_n(\gamma_1, \gamma_2)$ is convex in (γ_1, γ_2) and concave in θ while (67) follows since, as previously proven, under the condition $\delta < \delta_\star$, the set

$$\{(\gamma_1, \gamma_2) \mid (\gamma_1^2 + \gamma_2^2) \leq k^2 \text{ and } D_n(\gamma_1, \gamma_2) \leq 0\}$$

is not empty.

The function $(\gamma_1, \gamma_2) \mapsto D_n(\gamma_1, \gamma_2)$ is convex and converges pointwise to $(\gamma_1, \gamma_2) \mapsto \overline{D}(\gamma_1, \gamma_2, 0)$. Hence, it converges uniformly over compacts. For $\kappa > 0$, any $r \in \mathbb{N}^*$ and sufficiently large n , the following holds true:

$$\overline{D}(\gamma_1, \gamma_2) - \kappa \leq D_n(\gamma_1, \gamma_2) \leq \overline{D}(\gamma_1, \gamma_2) + \kappa, \quad \forall (\gamma_1, \gamma_2) \text{ such that } \gamma_1^2 + \gamma_2^2 \leq r^2$$

and thus:

$$\left\{ \{(\gamma_1, \gamma_2) \mid \overline{D}(\gamma_1, \gamma_2) \leq -\kappa \text{ and } \gamma_1^2 + \gamma_2^2 \leq r^2\} \right\} \subset \left\{ \{(\gamma_1, \gamma_2) \mid D_n(\gamma_1, \gamma_2) \leq 0\} \text{ and } \gamma_1^2 + \gamma_2^2 \leq r^2 \right\} \quad (68)$$

Before handling the proof of (64a) and (64b), we need first to establish that for integer k sufficiently large $\phi_k^{(n)}$ does not change with k , which will require often the use of (68). To start, for $\kappa > 0$, we consider the set:

$$\mathcal{D}_\kappa = \{(\gamma_1, \gamma_2) \mid \overline{D}(\gamma_1, \gamma_2) \leq -\kappa\}$$

We let $(\gamma_1^\kappa, \gamma_2^\kappa)$ be defined as:

$$(\gamma_1^\kappa, \gamma_2^\kappa) = \arg \min_{\substack{\gamma_1, \gamma_2 \\ \overline{D}(\gamma_1, \gamma_2) \leq -\kappa}} \frac{1}{2}(\gamma_2 - \beta)^2 + \frac{1}{2}\gamma_1^2$$

and define k_0 as $k_0 := \max \left(\lceil \sqrt{2} \left(\sqrt{(\gamma_2^\kappa - \beta)^2 + (\gamma_1^\kappa)^2} + \beta \right) \rceil + 1, 2\sqrt{(\gamma_1^\kappa)^2 + (\gamma_2^\kappa)^2} \right)$. Note that γ_1^κ and γ_2^κ are also given by:

$$(\gamma_1^\kappa, \gamma_2^\kappa) = \arg \min_{\substack{\gamma_1, \gamma_2 \\ \overline{D}(\gamma_1, \gamma_2) \leq -\kappa \\ \gamma_1^2 + \gamma_2^2 \leq k_0^2}} \frac{1}{2}(\gamma_2 - \beta)^2 + \frac{1}{2}\gamma_1^2$$

We will thus prove that for all $k \geq k_0$, with probability 1 for sufficiently large n ,

$$\phi_k^{(n)} = \phi_{k_0}^{(n)}. \quad (69)$$

To begin with, we invoke the uniform convergence of $(\gamma_1, \gamma_2) \mapsto \frac{1}{2}(\gamma_2 - \|\beta_\star\|)^2 + \frac{1}{2}\gamma_1^2$ to $(\gamma_1, \gamma_2) \mapsto \frac{1}{2}(\gamma_2 - \beta)^2 + \frac{1}{2}\gamma_1^2$, to claim that for any $\eta > 0$, and $(\gamma_1, \gamma_2) \in \{(\gamma_1, \gamma_2) \mid \gamma_1^2 + \gamma_2^2 \leq k_0^2\}$

$$\frac{1}{2}(\gamma_2 - \|\beta_\star\|)^2 + \frac{1}{2}\gamma_1^2 \leq \frac{1}{2}(\gamma_2 - \beta)^2 + \frac{1}{2}\gamma_1^2 + \eta$$

and hence,

$$\phi_{k_0}^{(n)} \leq \min_{D_n(\gamma_1, \gamma_2)} \frac{1}{2}(\gamma_2 - \beta)^2 + \frac{1}{2}\gamma_1^2 + \eta$$

Next, we use (68) to establish that

$$\phi_{k_0}^{(n)} \leq \frac{1}{2}(\gamma_2^\kappa - \beta)^2 + \frac{1}{2}(\gamma_1^\kappa)^2 + \eta. \quad (70)$$

Clearly, $\phi_k^{(n)} \leq \phi_{k_0}^{(n)}$. Let $\tilde{\gamma}_1$ and $\tilde{\gamma}_2$ be the minimizers of $\phi_k^{(n)}$. Then,

$$\phi_k^{(n)} < \phi_{k_0}^{(n)} \implies (\tilde{\gamma}_1, \tilde{\gamma}_2) \notin \{(\gamma_1, \gamma_2) \mid \gamma_1^2 + \gamma_2^2 \leq k_0^2\}$$

To prove (69), we proceed by contradiction and assume that

$$\phi_k^{(n)} < \phi_{k_0}^{(n)}. \quad (71)$$

Then, it is easy to see that in this case, we must have:

$$\tilde{\gamma}_1^2 + \tilde{\gamma}_2^2 \geq k_0^2$$

Therefore, either $\tilde{\gamma}_1^2 \geq \frac{1}{2}k_0^2$ or $\tilde{\gamma}_2^2 \geq \frac{1}{2}k_0^2$.

First case: $\tilde{\gamma}_1^2 \geq \frac{1}{2}k_0^2$. In this case:

$$\phi_k^{(n)} = \frac{1}{2}(\tilde{\gamma}_2 - \|\beta_\star\|)^2 + \frac{1}{2}\tilde{\gamma}_1^2 \geq \frac{1}{4}k_0^2 \geq \frac{1}{2}(\gamma_2^\kappa - \beta)^2 + \frac{1}{2}(\gamma_1^\kappa)^2 + \frac{1}{4} \geq \phi_{k_0}^{(n)}$$

where the last inequality follows from (70). This shows that $\phi_k^{(n)} \geq \phi_{k_0}^{(n)}$ which contradicts (71).

Second case: $\tilde{\gamma}_2^2 \geq \frac{1}{2}k_0^2$ Using the relation $(x - y)^2 \geq (|x| - |y|)^2$, we obtain

$$\phi_k^{(n)} \geq \frac{1}{2}(\sqrt{\frac{1}{2}k_0} - \|\beta_\star\|)^2 \geq \frac{1}{2} \left(\sqrt{(\gamma_2^\kappa - \beta)^2 + (\gamma_1^\kappa)^2} + \beta - \|\beta_\star\| + \frac{1}{\sqrt{2}} \right)^2 \geq \phi_{k_0}^{(n)}$$

where we used the fact $|\beta - \|\beta_\star\||$ can be assumed sufficiently small (let say smaller than $\frac{1}{2\sqrt{2}}$) for n sufficiently large. This again contradicts (71) which proves that $\phi_k^{(n)} = \phi_{k_0}^{(n)}$ for $k \geq k_0$ and n sufficiently large but independent of k . With the proof of (69) at hand, we are now ready to show (64a) and (64b).

Proof of (64a) The proof relies on using the uniform convergence of $(\gamma_1, \gamma_2) \mapsto \min_{\tilde{b}} D_n(\gamma_1, \gamma_2, \tilde{b})$ to $(\gamma_1, \gamma_2) \mapsto \overline{D}(\gamma_1, \gamma_2, 0)$ along with Lemma 4. More specifically, recalling (41), it holds that:

$$\left\{ \{(\gamma_1, \gamma_2) \mid \overline{D}_n(\gamma_1, \gamma_2) \leq 0 \text{ and } \gamma_1^2 + \gamma_2^2 \leq r^2\} \right\} \subset \left\{ \{(\gamma_1, \gamma_2) \mid \overline{D}(\gamma_1, \gamma_2) \leq \kappa\} \text{ and } \gamma_1^2 + \gamma_2^2 \leq r^2 \right\} \quad (72)$$

Hence,

$$\phi_{k_0}^{(n)} = \min_{\substack{\gamma_1^2 + \gamma_2^2 \leq k_0^2 \\ D_n(\gamma_1, \gamma_2) \leq 0}} \frac{1}{2}(\gamma_2 - \|\beta_\star\|)^2 + \frac{1}{2}\gamma_1^2 \quad (73)$$

$$\geq \min_{\substack{\gamma_1^2 + \gamma_2^2 \leq k_0^2 \\ \overline{D}(\gamma_1, \gamma_2) \leq \kappa}} \frac{1}{2}(\gamma_2 - \|\beta_\star\|)^2 + \frac{1}{2}\gamma_1^2 \quad (74)$$

Since $(\gamma_1, \gamma_2) \mapsto \frac{1}{2}(\gamma_2 - \|\beta_\star\|)^2 + \frac{1}{2}\gamma_1^2$ converges uniformly to $(\gamma_1, \gamma_2) \mapsto \frac{1}{2}(\gamma_2 - \beta)^2 + \frac{1}{2}\gamma_1^2$ over compacts, for any $\eta > 0$, there exists n sufficiently large such that:

$$\phi_{k_0}^{(n)} \geq \left(\min_{\substack{\gamma_1^2 + \gamma_2^2 \leq k_0^2 \\ \overline{D}(\gamma_1, \gamma_2) \leq \kappa}} \frac{1}{2}(\gamma_2 - \beta)^2 + \frac{1}{2}\gamma_1^2 \right) - \eta \geq \left(\min_{\substack{\gamma_1^2 + \gamma_2^2 \leq k_0^2 \\ \overline{D}(\gamma_1, \gamma_2) \leq \kappa}} \frac{1}{2}(\gamma_2 - \beta)^2 + \frac{1}{2}\gamma_1^2 \right) - 2\eta \quad (75)$$

where the last inequality follows by applying Lemma 4.

Proof of (64b) Fix any $\eta > 0$. It follows from (68) that

$$\phi_{k_0}^{(n)} \leq \left(\min_{\substack{\gamma_1^2 + \gamma_2^2 \leq k_0^2 \\ \overline{D}(\gamma_1, \gamma_2, 0) \leq -\kappa}} \frac{1}{2}(\gamma_2 - \beta)^2 + \frac{1}{2}\gamma_1^2 \right) + \eta \quad (76)$$

Similarly, we conclude by invoking Lemma 4 and using the η -definition of the infimum.

With (64a) and (64b) at hand, we now use Theorem 5 to conclude that:

$$\Phi^{(n)} = \|\hat{\mathbf{w}}_H\| \rightarrow \bar{\phi} = \frac{1}{2}(\gamma_2^* - \beta)^2 + \frac{1}{2}(\gamma_1^*)^2 \quad (77)$$

However, this convergence result is not sufficient to establish that of $\|\hat{\mathbf{w}}_H - \beta_\star\|^2$ and the cosine similarity. Hopefully, we can easily prove these results by working with the solution $\hat{\mathbf{w}}_H$ of (35) which, as per (36), writes as $\hat{\mathbf{w}}_H = -\hat{\mathbf{w}}_H + \beta_\star$. The norm of $\hat{\mathbf{w}}_H$ represents the risk while the cosine similarity can be retrieved by using the relation

$$\frac{\hat{\mathbf{w}}_H^T \beta_\star^T}{\|\beta_\star\|} = \frac{\hat{\mathbf{w}}_H^T \beta_\star}{\|\beta_\star\|} - \|\beta_\star\|$$

In the sequel, we will focus on the convergence of the quantity $\frac{\hat{\mathbf{w}}_H^T \beta_\star^T}{\|\beta_\star\|}$, from which the cosine similarity easily follows. The convergence of the risk can be done in a similar way and is omitted for brevity.

For that, we need to consider as suggested by Theorem 5 a perturbed version of the sequence of the AO problem in which $\tilde{\mathbf{w}}$ is further constrained on the set $\mathcal{S}_\xi$:

$$\mathcal{S}_\xi = \left\{ \tilde{\mathbf{w}} \in \mathbb{R}^p \mid \left| \frac{\tilde{\mathbf{w}}^T \beta_\star}{\|\beta_\star\|} - \gamma_2^* \right| > \xi \right\}$$

where ξ is any positive scalar. Particularly, we define for a given r and θ , $\tilde{\phi}_{r,\theta}$ as:

$$\tilde{\phi}_{r,\theta}^{(n)} = \min_{\substack{\|\tilde{\mathbf{w}}\| \leq r \\ \tilde{\mathbf{w}} \notin \mathcal{S}_\xi}} \max_{\|\mathbf{u}\| \leq \theta} \frac{1}{\sqrt{p}} \|\mathbf{P} \perp \tilde{\mathbf{w}}\|_2 \mathbf{g}^T \mathbf{u} + \frac{1}{\sqrt{p}} \|\mathbf{u}\| \mathbf{h}^T \mathbf{P} \perp \tilde{\mathbf{w}} + \frac{1}{\sqrt{p}} \frac{1}{\|\beta_\star\|} \mathbf{u}^T \mathbf{z} \beta_\star^T \tilde{\mathbf{w}} + \frac{1}{2} \|\beta_\star - \tilde{\mathbf{w}}\|^2 + \frac{1}{\sqrt{p}} \sigma \mathbf{n}^T \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T |\mathbf{u}|$$

Also, we define for $r \geq 0$, $\tilde{\phi}_r$ as:

$$\tilde{\phi}_r^{(n)} = \sup_{\theta \geq 0} \tilde{\phi}_{r,\theta}$$

Following the same calculations that led to (40), we can simplify $\tilde{\phi}_{r,\theta}^{(n)}$ as:

$$\tilde{\phi}_{r,\theta}^{(n)} = \min_{\substack{\gamma_1^2 + \gamma_2^2 \leq r^2 \\ |\gamma_2 - \gamma_2^*| \geq \xi}} \frac{1}{2}\gamma_1^2 + \frac{1}{2}\gamma_2^2 + \frac{1}{2}\|\beta_\star\|_2^2 - \gamma_2\|\beta_\star\|_2 + \theta(D_n(\gamma_1, \gamma_2))_+$$

Since the objective is convex in (γ_1, γ_2) , using (Deng et al., 2020, Remark 3), $\tilde{\phi}_r^{(n)}$ can be simplified as:

$$\tilde{\phi}_r^{(n)} = \min_{\substack{\gamma_1^2 + \gamma_2^2 \leq r^2 \\ |\gamma_2 - \gamma_2^*| \geq \xi}} \sup_{\theta \geq 0} \frac{1}{2}\gamma_1^2 + \frac{1}{2}\gamma_2^2 + \frac{1}{2}\|\beta_\star\|_2^2 - \gamma_2\|\beta_\star\|_2 + \theta(D_n(\gamma_1, \gamma_2))_+ = \min_{\substack{\gamma_1^2 + \gamma_2^2 \leq r^2 \\ |\gamma_2 - \gamma_2^*| \geq \xi \\ D_n(\gamma_1, \gamma_2) \leq 0}} \frac{1}{2}\gamma_1^2 + \frac{1}{2}\gamma_2^2 + \frac{1}{2}\|\beta_\star\|_2^2 - \gamma_2\|\beta_\star\|_2$$

By reference to Theorem 5, it suffices to prove that there exists $\zeta > 0$ such that for sufficiently large n (independent of r) it holds that:

$$\tilde{\phi}_r^{(n)} \geq \bar{\phi} + \zeta \quad (78)$$

or equivalently,

$$\forall r, \quad \tilde{\phi}_r^{(n)} \geq \bar{\phi} + \zeta \quad (79)$$

For that we proceed in two steps. In the first, step, we prove that for any $\tilde{\zeta} > 0$:

$$\tilde{\phi}_r^{(n)} \geq \tilde{\phi} - \tilde{\zeta} \quad (80)$$

where

$$\tilde{\phi} := \min_{\substack{\gamma_1, \gamma_2 \\ |\gamma_2 - \gamma_1^*| \geq \xi \\ \overline{D}(\gamma_1, \gamma_2) \leq 0}} \frac{1}{2}\gamma_1^2 + \frac{1}{2}\gamma_2^2 + \frac{1}{2}\beta^2 - \gamma_2\beta_2$$

The proof of (80) relies on the uniform convergence of D_n to $\overline{D}$ and is identical to what is done in (52). Details are thus omitted for brevity. In the second step, we show that there exists $\zeta > 0$ such that:

$$\tilde{\phi} \geq \overline{\phi} + 2\zeta. \quad (81)$$

To see this, it suffices to note that $\tilde{\phi} - \overline{\phi} > 0$. Indeed, clearly $\tilde{\phi} \geq \overline{\phi}$, however, since (γ_1^*, γ_2^*) are the unique minimizers of $\overline{\phi}$, we must have $\tilde{\phi} - \overline{\phi} > 0$. Hence, (81) holds for $\zeta = \frac{\tilde{\phi} - \overline{\phi}}{2}$. Starting from (80) and using $\tilde{\zeta}$ smaller than ζ , we obtain (78). Based on Theorem 5, we may conclude that the optimizer $\hat{\mathbf{w}}_H$ of (35) satisfies:

$$\frac{\hat{\mathbf{w}}_H^T \boldsymbol{\beta}_\star}{\|\boldsymbol{\beta}_\star\|} \xrightarrow{a.s.} \gamma_2^* \quad (82)$$

Hence,

$$\frac{\hat{\mathbf{w}}_H^T \boldsymbol{\beta}_\star}{\|\boldsymbol{\beta}_\star\|} \xrightarrow{a.s.} -\gamma_2^* + \beta_\star$$

Using (77), we thus obtain:

$$\frac{\hat{\mathbf{w}}_H^T \boldsymbol{\beta}_\star}{\|\boldsymbol{\beta}_\star\| \|\hat{\mathbf{w}}_H\|} \xrightarrow{a.s.} \frac{-\gamma_2^* + \beta_\star}{\sqrt{\frac{1}{2}(\gamma_2^* - \beta)^2 + \frac{1}{2}(\gamma_1^*)^2}}$$

Similarly, by defining the perturbed AO problem in which $\tilde{\mathbf{w}}$ is further constrained on the set

$$\mathbf{w} \in \left\{ \tilde{\mathbf{w}} \in \mathbb{R}^p \mid \left| \|\tilde{\mathbf{w}}\| - \sqrt{(\gamma_1^*)^2 + (\gamma_2^*)^2} \right| > \xi \right\}$$

and following the same approach, we can prove the convergence of the risk to $(\gamma_1^*)^2 + (\gamma_2^*)^2$. The results of Theorem 2 are then obtained by performing the change of variable $\tilde{\gamma}_1 = \frac{\gamma_1}{\sigma}$ and $\tilde{\gamma}_2 := \frac{\gamma_2}{\sigma}$.

8. Performance of S-SVR: Proof of Theorem 3

This section is devoted to the analysis of the statistical behavior of the S-SVR. To begin with, we shall note that the S-SVR can also be written as:

$$\begin{aligned} \min_{\mathbf{w}} \quad & \frac{1}{2} \|\mathbf{w}\|^2 + \frac{C}{p} \sum_{i=1}^n \xi_i \\ \text{s.t.} \quad & y_i - \mathbf{w}^T \mathbf{x}_i \leq \epsilon + \xi_i, \quad i = 1, \dots, n \\ & \mathbf{w}^T \mathbf{x}_i - y_i \leq \epsilon + \xi_i, \\ & \xi_i \geq 0, \end{aligned} \quad (83)$$

where in (83) we replaced $\tilde{\xi}_i$ by ξ_i . The reason why (3) and (83) are equivalent is as follows: First it is easy to see that (83) is identical to (3) when ξ_i and $\tilde{\xi}_i$ are constrained to be equal. Hence, for any $(\mathbf{w}, \{\xi_i\}_{i=1}^n)$ feasible for (83), $(\mathbf{w}, \{\xi_i\}_{i=1}^n, \{\xi_i\}_{i=1}^n)$ is also feasible for (3). Next, let $(\mathbf{w}, \{\xi_i\}_{i=1}^n, \{\tilde{\xi}_i\}_{i=1}^n)$ be feasible for (3). Then, for all $i = 1, \dots, n$,

$$\begin{aligned} y_i - \mathbf{w}^T \mathbf{x}_i \leq \epsilon + \xi_i \text{ and } \mathbf{w}^T \mathbf{x}_i - y_i \leq \epsilon + \tilde{\xi}_i &\implies |\mathbf{w}^T \mathbf{x}_i - y_i| \leq \epsilon + \max(\xi_i, \tilde{\xi}_i) \\ &\implies \mathbf{w}^T \mathbf{x}_i - y_i \leq \epsilon + \max(\xi_i, \tilde{\xi}_i) \text{ and } y_i - \mathbf{w}^T \mathbf{x}_i \leq \epsilon + \max(\xi_i, \tilde{\xi}_i) \end{aligned}$$

From this it follows that if $(\mathbf{w}, \{\xi_i\}_{i=1}^n, \{\tilde{\xi}_i\}_{i=1}^n)$ is feasible for (3), then $(\mathbf{w}, \{\max(\xi_i, \tilde{\xi}_i)\}_{i=1}^n)$ is feasible for (83). The optimal costs for (3) and (83) are as such identical and thus are their respective solutions due to the strict convexity of their objectives. In the sequel, for the sake of simplicity, we will study (83) instead of (3).

Identifying the PO. The Lagrangian associated with problem (83) can be written as:

$$\mathcal{L}(\mathbf{w}, \boldsymbol{\xi}, \boldsymbol{\lambda}, \boldsymbol{\alpha}) = \frac{1}{2} \|\mathbf{w}\|^2 + \frac{C}{p} \sum_{i=1}^n \xi_i + \sum_{i=1}^n \lambda_i (\beta_\star^T \mathbf{x}_i - \mathbf{w}^T \mathbf{x}_i + \sigma n_i - \epsilon - \xi_i) + \sum_{i=1}^n \alpha_i (\mathbf{w}^T \mathbf{x}_i - \beta_\star^T \mathbf{x}_i - \sigma n_i - \epsilon - \xi_i)$$

where $\boldsymbol{\xi} = [\xi_1, \dots, \xi_n]$ and $\boldsymbol{\lambda} = [\lambda_1, \dots, \lambda_n]$ and $\boldsymbol{\alpha} = [\alpha_1, \dots, \alpha_n]$ are the Lagrangian coefficients. Letting $\tilde{\mathbf{w}} = \beta_\star - \mathbf{w}$, we have

$$\mathcal{L}(\tilde{\mathbf{w}}, \boldsymbol{\xi}, \boldsymbol{\lambda}, \boldsymbol{\alpha}) = \frac{1}{2} \|\beta_\star - \tilde{\mathbf{w}}\|^2 + \frac{C}{p} \sum_{i=1}^n \xi_i + \sum_{i=1}^n \lambda_i (\tilde{\mathbf{w}}^T \mathbf{x}_i + \sigma n_i - \epsilon - \xi_i) + \sum_{i=1}^n \alpha_i (-\tilde{\mathbf{w}}^T \mathbf{x}_i - \sigma n_i - \epsilon - \xi_i)$$

Writing $\mathbf{X} = [\mathbf{x}_1, \dots, \mathbf{x}_n]$ leads to the following optimization problem:

$$\min_{\tilde{\mathbf{w}}, \boldsymbol{\xi} \geq 0} \max_{\boldsymbol{\lambda} \geq 0, \boldsymbol{\alpha} \geq 0} \frac{1}{2} \|\beta_\star - \tilde{\mathbf{w}}\|^2 + \sum_{i=1}^n \left(\frac{C}{p} - \lambda_i - \alpha_i \right) \xi_i + \tilde{\mathbf{w}}^T \mathbf{X}(\boldsymbol{\lambda} - \boldsymbol{\alpha}) + \sigma \mathbf{n}^T (\boldsymbol{\lambda} - \boldsymbol{\alpha}) - \epsilon \mathbf{1}^T (\boldsymbol{\lambda} + \boldsymbol{\alpha}) \quad (84)$$

Let $\boldsymbol{\xi}^\star$ be the optimum in $\boldsymbol{\xi}$ of the above problem. It follows from the first order condition that for all $\boldsymbol{\xi} \geq 0$, the following inequality must hold

$$\left(\frac{C}{p} - \lambda_i - \alpha_i \right) (\xi_i - \xi_i^\star) \geq 0, \quad i = 1, \dots, n \quad (85)$$

For (85) to be satisfied, we must have $\frac{C}{p} - \lambda_i - \alpha_i \geq 0$ and $\left(\frac{C}{p} - \lambda_i - \alpha_i \right) \xi_i^\star = 0$. Using this, the problem in (84) simplifies as:

$$\min_{\tilde{\mathbf{w}}} \max_{\substack{\boldsymbol{\lambda} \geq 0, \boldsymbol{\alpha} \geq 0 \\ \boldsymbol{\lambda} + \boldsymbol{\alpha} \leq \frac{C}{p}}} \frac{1}{2} \|\beta_\star - \tilde{\mathbf{w}}\|^2 + \tilde{\mathbf{w}}^T \mathbf{X}(\boldsymbol{\lambda} - \boldsymbol{\alpha}) + \sigma \mathbf{n}^T (\boldsymbol{\lambda} - \boldsymbol{\alpha}) - \epsilon \mathbf{1}^T (\boldsymbol{\lambda} + \boldsymbol{\alpha})$$

Considering $\mathbf{u} = \sqrt{p}(\boldsymbol{\lambda} - \boldsymbol{\alpha})$ and $\mathbf{v} = \sqrt{p}(\boldsymbol{\lambda} + \boldsymbol{\alpha})$. With these notations, the optimization problem can be written as

$$\min_{\tilde{\mathbf{w}}} \max_{\substack{\mathbf{u} \geq 0 \\ |\mathbf{u}| \leq \mathbf{v} \\ \mathbf{v} \leq \frac{C}{\sqrt{p}}}} \frac{1}{2} \|\beta_\star - \tilde{\mathbf{w}}\|^2 + \frac{1}{\sqrt{p}} \tilde{\mathbf{w}}^T \mathbf{X} \mathbf{u} + \frac{\sigma}{\sqrt{p}} \mathbf{n}^T \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T \mathbf{v} \quad (86)$$

Clearly, and for the same arguments as in the proof of the H-SVR, the optimization over $\mathbf{v}$ leads to $\mathbf{v}^\star = |\mathbf{u}|$. Replacing $\mathbf{v}$ by its optimum value, (86) simplifies to,

$$\tilde{\Phi} := \min_{\tilde{\mathbf{w}}} \max_{\substack{\mathbf{u} \geq 0 \\ |\mathbf{u}| \leq \frac{C}{\sqrt{p}}}} \frac{1}{2} \|\beta_\star - \tilde{\mathbf{w}}\|^2 + \frac{1}{\sqrt{p}} \tilde{\mathbf{w}}^T \mathbf{X} \mathbf{u} + \frac{\sigma}{\sqrt{p}} \mathbf{n}^T \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T |\mathbf{u}| \quad (87)$$

Contrary to the H-SVR, the S-SVR is always feasible. As such, the optimal cost is finite and thus should be attained by a solution with finite norm. As far as the proof is concerned, checking that at optimum, the solution has a bounded norm allows us to work with the original CGMT framework in (Thrampoulidis et al., 2018). Whenever this possible, proceeding in this way should be preferred for the sake of simplicity. We thus start by proving that the norm of the optimum solution $\tilde{\mathbf{w}}^\star$ is bounded. Indeed, it follows from the first order optimality conditions that:

$$\mathbf{w}^\star = \beta_\star - \frac{1}{\sqrt{p}} \mathbf{X} \mathbf{u}$$

and hence,

$$\|\tilde{\mathbf{w}}^\star\| \leq \|\beta_\star\| + \frac{1}{2} \left\| \frac{1}{\sqrt{p}} \mathbf{X} \right\| \|\mathbf{u}\|$$

Clearly, the constraint $|\mathbf{u}| \leq \frac{C}{\sqrt{p}}$ implies that $\|\mathbf{u}\|$ is bounded. Moreover, it is known from standard results of random matrix theory that $\left\| \frac{1}{\sqrt{p}} \mathbf{X} \right\|$ is almost surely bounded (Bai & Silverstein, 2009) and hence $\tilde{\mathbf{w}}$ can be assumed without loss of

generally to satisfy $\|\tilde{\mathbf{w}}\| \leq C_w$ where C_w is a finite constant. With this, the optimization problem can be written in the form of a primary optimization problem as required by the CGMT framework in (Thrapoulidis et al., 2018):

$$\Phi_S^{(n)} = \min_{\|\tilde{\mathbf{w}}\| \leq C_w} \max_{\mathbf{u} \leq \frac{C}{\sqrt{p}}} \frac{1}{\sqrt{p}} \tilde{\mathbf{w}}^T \mathbf{X} \mathbf{u} + \frac{\sigma}{\sqrt{p}} \mathbf{n}^T \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T |\mathbf{u}| + \frac{1}{2} \|\beta_\star - \tilde{\mathbf{w}}\|^2. \quad (88)$$

To connect (87) to (88), we rely on (Thrapoulidis et al., 2018, Lemma 5), which ensures that if there exists $\gamma_1^\star$ and $\gamma_2^\star$ such that for sufficiently large C_w , the optimizer $\tilde{\mathbf{w}}_{C_w}$ satisfies:

$$\|\tilde{\mathbf{w}}_{C_w}\| \xrightarrow{a.s.} (\gamma_1^\star)^2 + (\gamma_2^\star)^2, \quad \frac{\tilde{\mathbf{w}}_{C_w}^T \beta_\star}{\|\beta_\star\|} \xrightarrow{a.s.} \gamma_2^\star \quad (89)$$

then any minimizer $\tilde{\mathbf{w}}$ of (87) satisfies:

$$\|\tilde{\mathbf{w}}\| \xrightarrow{a.s.} (\gamma_1^\star)^2 + (\gamma_2^\star)^2, \quad \frac{\tilde{\mathbf{w}}^T \beta_\star}{\|\beta_\star\|} \xrightarrow{a.s.} \gamma_2^\star \quad (90)$$

From now onward, we thus focus on analyzing $\Phi_S^{(n)}$. Using the same trick as in the proof for the H-SVR, we may project $\tilde{\mathbf{w}}$ onto $\beta_\star$ and the space orthogonal to it, thereby yielding:

$$\Phi_S^{(n)} = \min_{\|\tilde{\mathbf{w}}\| \leq C_w} \max_{\mathbf{u} \leq \frac{C}{\sqrt{p}}} \frac{1}{\sqrt{p}} \tilde{\mathbf{w}}^T \mathbf{P}_\perp \mathbf{X} \mathbf{u} + \frac{1}{\sqrt{p} \|\beta_\star\|^2} \tilde{\mathbf{w}}^T \beta_\star \beta_\star^T \mathbf{X} \mathbf{u} + \frac{\sigma}{\sqrt{p}} \mathbf{n}^T \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T |\mathbf{u}| + \frac{1}{2} \|\beta_\star - \tilde{\mathbf{w}}\|^2. \quad (91)$$

where we recall that $\mathbf{P}_\perp = \mathbf{I}_p - \frac{\beta_\star \beta_\star^T}{\beta_\star^T \beta_\star}$. Let $\mathbf{z} = \frac{1}{\|\beta_\star\|} \mathbf{X}^T \beta_\star$. Then, $\mathbf{z}$ is independent of $\mathbf{P}_\perp \mathbf{X}$ and $\Phi_S^{(n)}$ writes as:

$$\Phi_S^{(n)} = \min_{\|\tilde{\mathbf{w}}\| \leq C_w} \max_{\mathbf{u} \leq \frac{C}{\sqrt{p}}} \frac{1}{\sqrt{p}} \tilde{\mathbf{w}}^T \mathbf{P}_\perp \mathbf{X} \mathbf{u} + \frac{1}{\sqrt{p} \|\beta_\star\|} \tilde{\mathbf{w}}^T \beta_\star \mathbf{z}^T \mathbf{u} + \frac{\sigma}{\sqrt{p}} \mathbf{n}^T \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T |\mathbf{u}| + \frac{1}{2} \|\beta_\star - \tilde{\mathbf{w}}\|^2. \quad (92)$$

Identification and simplification of the AO. With the primary problem in (92), we associate the following AO problem given by:

$$\phi_S^{(n)} = \min_{\|\tilde{\mathbf{w}}\| \leq C_w} \max_{\mathbf{u} \leq \frac{C}{\sqrt{p}}} \frac{1}{\sqrt{p}} \|\mathbf{P}_\perp \tilde{\mathbf{w}}\| \mathbf{g}^T \mathbf{u} - \frac{1}{\sqrt{p}} \|\mathbf{u}\| \mathbf{h}^T \mathbf{P}_\perp \tilde{\mathbf{w}} + \frac{1}{\sqrt{p} \|\beta_\star\|} \tilde{\mathbf{w}}^T \beta_\star \mathbf{z}^T \mathbf{u} + \frac{\sigma}{\sqrt{p}} \mathbf{n}^T \mathbf{u} - \frac{1}{\sqrt{p}} \epsilon \mathbf{1}^T |\mathbf{u}| + \frac{1}{2} \|\beta_\star - \tilde{\mathbf{w}}\|^2 \quad (93)$$

In a similar way as in the H-SVR, we decompose $\tilde{\mathbf{w}}$ as:

$$\tilde{\mathbf{w}} = \gamma_1 \frac{\mathbf{P}_\perp \tilde{\mathbf{w}}}{\|\mathbf{P}_\perp \tilde{\mathbf{w}}\|} + \gamma_2 \frac{\beta_\star}{\|\beta_\star\|}$$

Hence,

$$\phi_S^{(n)} = \min_{\gamma_1, \gamma_2} \max_{\mathbf{u} \leq \frac{C}{\sqrt{p}}} A_n(\gamma_1, \gamma_2, \mathbf{u}) \quad (94)$$

where

$$A_n(\gamma_1, \gamma_2, \mathbf{u}) = \frac{\gamma_1}{\sqrt{p}} \mathbf{g}^T \mathbf{u} - \frac{\gamma_1}{\sqrt{p}} \|\mathbf{u}\| \|\mathbf{P}_\perp \mathbf{h}\| + \frac{\gamma_2}{\sqrt{p}} \mathbf{z}^T \mathbf{u} + \frac{\sigma}{\sqrt{p}} \mathbf{n}^T \mathbf{u} - \frac{\epsilon}{\sqrt{p}} \|\mathbf{u}\|_1 + \frac{1}{2} (\gamma_1^2 + \gamma_2^2) + \frac{1}{2} \|\beta_\star\|^2 - \gamma_1 \|\beta_\star\|$$

It can be easily seen that

$$\inf_{\gamma_1, \gamma_2} \sup_{\mathbf{u} \leq \frac{C}{\sqrt{p}}} |A_n(\gamma_1, \gamma_2, \mathbf{u}) - \tilde{A}_n(\gamma_1, \gamma_2, \mathbf{u})| \xrightarrow{a.s.} 0.$$

where

$$\tilde{A}_n(\gamma_1, \gamma_2, \mathbf{u}) := \frac{\gamma_1}{\sqrt{p}} \mathbf{g}^T \mathbf{u} - \gamma_1 \|\mathbf{u}\| + \frac{\gamma_2}{\sqrt{p}} \mathbf{z}^T \mathbf{u} + \frac{\sigma}{\sqrt{p}} \mathbf{n}^T \mathbf{u} - \frac{\epsilon}{\sqrt{p}} \|\mathbf{u}\|_1 + \frac{1}{2} (\gamma_1^2 + \gamma_2^2) + \frac{1}{2} \beta^2 - \gamma_1 \beta$$

As a result $\phi_S^{(n)} - \tilde{\phi}_S^{(n)} \xrightarrow{a.s.} 0$, where

$$\tilde{\phi}_S^{(n)} = \min_{\substack{\gamma_1, \gamma_2 \\ \gamma_1^2 + \gamma_2^2 \leq C_w^2}} \max_{\substack{\mathbf{u} \\ \|\mathbf{u}\| \leq \frac{C}{\sqrt{p}}}} \tilde{A}_n(\gamma_1, \gamma_2, \mathbf{u})$$

Based on this, we will work from now on with $\tilde{\phi}_S^{(n)}$. For that, we rely on Lemma 5 to perform the optimization with respect to $\mathbf{u}$ and simplify $\tilde{\phi}_S^{(n)}$ as:

$$\tilde{\phi}_S^{(n)} = \min_{\substack{\gamma_1, \gamma_2 \\ \gamma_1^2 + \gamma_2^2 \leq C_w^2}} \sup_{\chi \geq 0} \frac{1}{2}(\gamma_1^2 + \gamma_2^2) + \frac{1}{2}\beta^2 - \gamma_1\beta + \hat{R}_n(\gamma_1, \gamma_2, \chi) \quad (95)$$

where

$$\hat{R}_n(\gamma_1, \gamma_2, \chi) := \begin{cases} \frac{1}{p} \sum_{i=1}^n C(b_i - \frac{C\gamma_1}{2\chi}) \mathbf{1}_{\{b_i\chi > \gamma_1 C\}} + \frac{1}{p} \sum_{i=1}^n \frac{b_i^2 \chi}{2\gamma_1} \mathbf{1}_{\{b_i\chi \leq \gamma_1 C\}} - \frac{\gamma_1 \chi}{2}, & \text{if } \gamma_1 \neq 0 \\ \frac{C}{p} \sum_{i=1}^n b_i & \text{if } \gamma_1 = 0. \end{cases}$$

with $\mathbf{b} = (|\gamma_1 \mathbf{g} + \gamma_2 \mathbf{z} + \sigma \mathbf{n}| - \epsilon)_+$.

Asymptotic behavior of the AOs costs. Denote the objective function of (95) by $\hat{D}_n(\gamma_1, \gamma_2, \chi)$. Fix γ_1, γ_2, χ , then:

$$\hat{D}_n(\gamma_1, \gamma_2, \chi) \xrightarrow{a.s.} \bar{D}(\gamma_1, \gamma_2, \chi) := \frac{1}{2}\gamma_2^2 + \frac{1}{2}(\gamma_1 - \beta)^2 + \bar{R}(\gamma_1, \gamma_2, \chi) \quad (96)$$

where $\bar{R}(\gamma_1, \gamma_2, \chi) = \lim_{n \rightarrow \infty} \hat{R}_n(\gamma_1, \gamma_2, \chi)$ and is given by:

$$\bar{R}(\gamma_1, \gamma_2, \chi) := \begin{cases} \delta \mathbb{E} \left\{ C \left[\left(\left| \sqrt{\gamma_1^2 + \gamma_2^2} G + \sigma N \right| - \epsilon \right)_+ - \frac{C\gamma_1}{2\chi} \right] \mathbf{1}_{\{(|\sqrt{\gamma_1^2 + \gamma_2^2} G + \sigma N| - \epsilon)_+ \chi \geq \gamma_1 C\}} \right. \\ \left. + \frac{\chi}{2\gamma_1} \left(\left| \sqrt{\gamma_1^2 + \gamma_2^2} G + \sigma N \right| - \epsilon \right)_+^2 \mathbf{1}_{\{(|\sqrt{\gamma_1^2 + \gamma_2^2} G + \sigma N| - \epsilon)_+ \chi \leq \gamma_1 C\}} \right\} - \frac{\gamma_1 \chi}{2}, & \text{if } \gamma_1 \neq 0 \\ \tilde{R}(\gamma_1, \gamma_2) := C \delta \mathbb{E} [(|\gamma_2 G + \sigma N| - \epsilon)_+], & \text{if } \gamma_1 = 0. \end{cases}$$

Fix $\gamma_1 \neq 0$ and $\gamma_2 \in \mathbb{R}$. Function $(\gamma_1, \gamma_2) \mapsto \sup_{\chi \geq 0} \hat{D}_n(\gamma_1, \gamma_2, \chi)$ is convex in its arguments. Moreover, one can check that for $\gamma_1 \neq 0$:

$$\lim_{\chi \rightarrow \infty} \bar{D}(\gamma_1, \gamma_2, \chi) = -\infty$$

Hence, we may use (Thrapoulidis et al., 2018, Lemma 10) to obtain:

$$\sup_{\chi \geq 0} \hat{D}_n(\gamma_1, \gamma_2, \chi) \xrightarrow{a.s.} \sup_{\chi \geq 0} \bar{D}(\gamma_1, \gamma_2, \chi), \quad \gamma_1 \neq 0, \gamma_2 \in \mathbb{R} \quad (97)$$

When $\gamma_1 = 0$, clearly,

$$\sup_{\chi \geq 0} \hat{D}_n(0, \gamma_2, \chi) \xrightarrow{a.s.} \sup_{\chi \geq 0} \bar{D}(0, \gamma_2, \chi)$$

Function $(\gamma_1, \gamma_2) \mapsto \sup_{\chi \geq 0} \hat{D}_n(\gamma_1, \gamma_2, \chi)$ is convex in its arguments and converges pointwise to $(\gamma_1, \gamma_2) \mapsto \sup_{\chi \geq 0} \bar{D}(\gamma_1, \gamma_2, \chi)$. It thus converges uniformly over compacts in $\mathbb{R}^2$. Hence,

$$\min_{\substack{\gamma_1, \gamma_2 \\ \gamma_1^2 + \gamma_2^2 \leq C_w^2}} \sup_{\chi \geq 0} \hat{D}_n(\gamma_1, \gamma_2, \chi) \xrightarrow{a.s.} \min_{\substack{\gamma_1, \gamma_2 \\ \gamma_1^2 + \gamma_2^2 \leq C_w^2}} \sup_{\chi \geq 0} \bar{D}(\gamma_1, \gamma_2, \chi)$$

Now, one can easily check that

$$\lim_{\|\gamma_2\| \rightarrow \infty} \sup_{\chi \geq 0} \bar{D}(\gamma_1, \gamma_2, \chi) \geq \lim_{\|\gamma_2\| \rightarrow \infty} \frac{1}{2}(\gamma_1^2 + \gamma_2^2) - \gamma_1\beta = \infty$$

Hence,

$$\min_{\substack{\gamma_1, \gamma_2 \\ \gamma_1^2 + \gamma_2^2 \leq C_w^2}} \sup_{\chi \geq 0} \bar{D}(\gamma_1, \gamma_2, \chi) = \min_{\gamma_1, \gamma_2} \sup_{\chi \geq 0} \bar{D}(\gamma_1, \gamma_2, \chi)$$

and as such:

$$\min_{\substack{\gamma_1, \gamma_2 \\ \gamma_1^2 + \gamma_2^2 \leq C_w^2}} \sup_{\chi \geq 0} \hat{D}_n(\gamma_1, \gamma_2, \chi) \xrightarrow{a.s.} \min_{\gamma_1, \gamma_2} \sup_{\chi \geq 0} \bar{D}(\gamma_1, \gamma_2, \chi)$$

Concluding. Applying the CGMT framework, we conclude that:

$$\Phi_S^{(n)} \xrightarrow{a.s.} \min_{\gamma_1, \gamma_2} \sup_{\chi \geq 0} \bar{D}(\gamma_1, \gamma_2, \chi)$$

Since $(\gamma_1, \gamma_2) \mapsto \sup_{\chi \geq 0} \bar{D}(\gamma_1, \gamma_2, \chi)$ is jointly strictly convex in its arguments and is coercive, it admits a unique minimizer (γ_1^*, γ_2^*) . Considering a suitable perturbation of the AO problem and applying the same arguments as in H-SVR, we may conclude (89) and (90). Similarly to the H-SVR, to retrieve the results of Theorem 3, we use the change of variable $\tilde{\gamma}_1 = \frac{\gamma_1}{\sigma}$ and $\tilde{\gamma}_2 = \frac{\gamma_2}{\sigma}$.

9. Technical Lemmas

Lemma 1. *Let m be a strictly positive scalar and $\mathbf{a}$ be a vector in $\mathbb{R}^n$. Then:*

$$\max_{\substack{\mathbf{u} \in \mathbb{R}^n \\ \|\mathbf{u}\|_2 = m}} \mathbf{u}^T \mathbf{a} - \epsilon \|\mathbf{u}\|_1 = m \sqrt{\sum_{i=1}^n (|a_i| - \epsilon)_+^2} \quad (98)$$

Proof. To begin with, we note that if for some i , the optimal solution is non-zero then it should have the same sign as a_i . Hence, (98) is equivalent to the following optimization problem:

$$\max_{\substack{\mathbf{u} \in \mathbb{R}^n \\ \|\mathbf{u}\|_2 = m}} \sum_{i=1}^n |u_i| (|a_i| - \epsilon) \quad (99)$$

If for some i , $|a_i| \leq \epsilon$ then we need to set $u_i = 0$. Hence, Problem (99) amounts to solving:

$$\max_{\substack{\mathbf{u} \in \mathbb{R}^n \\ \|\mathbf{u}\|_2 = m}} \sum_{i=1}^n |u_i| (|a_i| - \epsilon)_+$$

Using Cauchy-Schwartz inequality, the objective of the above optimization problem can be upper-bounded as

$$\sum_{i=1}^n |u_i| (|a_i| - \epsilon)_+ \leq \|\mathbf{u}\|_2 \sqrt{\sum_{i=1}^n (|a_i| - \epsilon)_+^2}$$

where equality holds when $\mathbf{u} = m \frac{(|\mathbf{a}| - \epsilon)_+}{\sqrt{\|(|\mathbf{a}| - \epsilon)_+\|^2}}$. The optimal cost of (98) is thus given by $m \sqrt{\sum_{i=1}^n (|a_i| - \epsilon)_+^2}$. $\square$

Lemma 2. *Let $f : \mathbb{R} \mapsto \mathbb{R}$ be a convex and even function. Then for all $x \in \mathbb{R}$,*

$$f(0) \leq f(x)$$

or in other words f is minimized at zero.

Proof. Let $x \in \mathbb{R}$. Then, by convexity of f ,

$$f(0) = f\left(\frac{1}{2}x - \frac{1}{2}x\right) \leq \frac{1}{2}f(x) + \frac{1}{2}f(-x)$$

As $f(x) = f(-x)$, we thus have:

$$f(0) \leq f(x).$$

$\square$

Lemma 3. (Boyd & Vandenberghe, 2003) Let X and Y be two convex sets. Let $f : X \times Y \rightarrow \mathbb{R}$ be a jointly convex function in $X \times Y$. Assume that $\forall y \in Y, \inf_{x \in X} f(x, y) > -\infty$. Then, $g : y \rightarrow \inf_{x \in X} f(x, y)$ is convex in Y .

Lemma 4. (Kammoun & Alouini, 2020a) Let $d \in \mathbb{N}^*$. Let S_x be a compact non-empty set in $\mathbb{R}^d$. Let f and c be two continuous functions over S_x such that the set $\{c(x) \leq 0\}$ is non-empty. Then:

$$\min_{\substack{\mathbf{x} \in S_x \\ c(\mathbf{x}) \leq 0}} f(\mathbf{x}) = \sup_{\delta \geq 0} \min_{\substack{\mathbf{x} \in S_x \\ c(\mathbf{x}) \leq \delta}} f(\mathbf{x}) = \inf_{\delta > 0} \min_{\substack{\mathbf{x} \in S_x \\ c(\mathbf{x}) \leq -\delta}} f(\mathbf{x})$$

Lemma 5. Let $\mathbf{a} \in \mathbb{R}^{n \times 1}$. Let β, ϵ , and τ be positive scalars. Then,

- if $\beta = 0$,

$$\max_{|\mathbf{u}| \leq \tau} (\mathbf{a} - \epsilon \text{sign}(\mathbf{u}))^T \mathbf{u} - \beta \|\mathbf{u}\| = \max_{|\mathbf{u}| \leq \tau} (\mathbf{a} - \epsilon \text{sign}(\mathbf{u}))^T \mathbf{u} = \sum_{i=1}^n \tau \max(|a_i| - \epsilon, 0), \quad (100)$$

- if $\beta \neq 0$,

$$\max_{|\mathbf{u}| \leq \tau} (\mathbf{a} - \epsilon \text{sign}(\mathbf{u}))^T \mathbf{u} - \beta \|\mathbf{u}\| = \sup_{\chi > 0} \sum_{i=1}^n \frac{b_i^2 \chi}{2\beta} \mathbb{1}_{\{\frac{b_i \chi}{\beta} \leq \tau\}} + \sum_{i=1}^n \left(b_i \tau - \frac{\beta}{2\chi} \tau^2 \right) \mathbb{1}_{\{\frac{b_i \chi}{\beta} > \tau\}} - \beta \frac{\chi}{2} \quad (101)$$

where $b_i = (|a_i| - \epsilon)_+$. Moreover, function $\chi \mapsto \sum_{i=1}^n \frac{b_i^2 \chi}{2\beta} \mathbb{1}_{\{\frac{b_i \chi}{\beta} \leq \tau\}} + \sum_{i=1}^n \left(b_i \tau - \frac{\beta}{2\chi} \tau^2 \right) \mathbb{1}_{\{\frac{b_i \chi}{\beta} > \tau\}} - \beta \frac{\chi}{2}$ is concave on $(0, +\infty)$.

Proof. If $\beta = 0$, one can note that if for some i the optimal solution is non zero then it should have the same sign as a_i . Hence, the optimization problem becomes:

$$\min_{\substack{0 \leq |u_i| \leq \tau \\ i=1, \dots, n}} \sum_{i=1}^n |u_i| (|a_i| - \epsilon) = \sum_{i=1}^n \tau (|a_i| - \epsilon)_+$$

To treat the case $\beta \neq 0$, we start by rewriting $\|\mathbf{u}\|$ as

$$\|\mathbf{u}\| = \inf_{\chi > 0} \frac{\chi}{2} + \frac{\|\mathbf{u}\|^2}{2\chi},$$

thus yielding:

$$\max_{|\mathbf{u}| \leq \tau} (\mathbf{a} - \epsilon \text{sign}(\mathbf{u}))^T \mathbf{u} - \beta \|\mathbf{u}\| = \max_{|\mathbf{u}| \leq \tau} \sup_{\chi > 0} (\mathbf{a} - \epsilon \text{sign}(\mathbf{u}))^T \mathbf{u} - \beta \left[\frac{\chi}{2} + \frac{\|\mathbf{u}\|^2}{2\chi} \right]$$

In a similar way as in the case of $\beta = 0$, we note that if for some i the optimum u_i is non zero then it necessarily have the same sign as a_i . Moreover, if for some i , $|a_i| \leq \epsilon$ then necessarily $u_i = 0$. With these, the above problem writes thus as:

$$\sup_{\chi > 0} \max_{|\mathbf{u}| \leq \tau} \sum_{i=1}^n (|a_i| - \epsilon)_+ |u_i| - \frac{\beta}{2\chi} u_i^2 - \beta \frac{\chi}{2}.$$

Let $\mathbf{v} = |\mathbf{u}|$, and define for $i = 1, \dots, n$ $b_i = (|a_i| - \epsilon)_+$. Now, noting that $b_i \geq 0$ and that the function $x \rightarrow b_i x - \frac{\beta}{2\chi} x^2$ is increasing on $(-\infty, \frac{b_i \chi}{\beta})$ and decreasing on $(\frac{b_i \chi}{\beta}, \infty)$ with its maximum achieved at $x^* = \frac{b_i \chi}{\beta}$, one can easily see that

$$\max_{0 \leq x \leq \tau} b_i x - \frac{\beta}{2\chi} x^2 = \begin{cases} b_i \tau - \frac{\beta}{2\chi} \tau^2 & \text{if } \frac{b_i \chi}{\beta} > \tau \\ \frac{b_i^2 \chi}{2\beta} & \text{if } \frac{b_i \chi}{\beta} \leq \tau. \end{cases}$$

Hence,

$$\begin{aligned} & \sup_{\chi > 0} \max_{|\mathbf{u}| \leq \tau} \sum_{i=1}^n (a_i - \epsilon \text{sign}(u_i)) u_i - \frac{\beta}{2\chi} u_i^2 - \beta \frac{\chi}{2} \\ &= \sup_{\chi > 0} \sum_{i=1}^n \frac{b_i^2 \chi}{2\beta} \mathbb{1}_{\{\frac{b_i \chi}{\beta} \leq \tau\}} + \sum_{i=1}^n \left(b_i \tau - \frac{\beta}{2\chi} \tau^2 \right) \mathbb{1}_{\{\frac{b_i \chi}{\beta} > \tau\}} - \beta \frac{\chi}{2} \end{aligned}$$

Now, we proceed with the proof of the concavity of the function

$$\psi : \chi \rightarrow \sum_{i=1}^n \frac{b_i^2 \chi}{2\beta} \mathbb{1}_{\{\frac{b_i \chi}{\beta} \leq \tau\}} + \sum_{i=1}^n \left(b_i \tau - \frac{\beta}{2\chi} \tau^2 \right) \mathbb{1}_{\{\frac{b_i \chi}{\beta} > \tau\}} - \beta \frac{\chi}{2}.$$

For that, it suffices to note that

$$\psi(\chi) = \max_{|u_i| \leq \tau} \sum_{i=1}^n a_i u_i - \epsilon |u_i| - \frac{\beta}{2\chi} u_i^2 - \frac{\beta \chi}{2}$$

Function $(u_i, \chi) \mapsto \frac{\beta u_i^2}{2\chi}$ is jointly convex on $[-\tau, \tau] \times \mathbb{R}_+$, then $(u_i, \chi) \mapsto a_i u_i - \epsilon |u_i| - \frac{\beta}{2\chi} u_i^2 - \frac{\beta \chi}{2}$ is jointly concave on $[-\tau, \tau] \times \mathbb{R}_+$. Applying Lemma 3, we show that function ψ is concave. $\square$