

For interpolating kernel machines, minimizing the norm of the ERM solution maximizes stability

Akshay Rangamani^{*,†}, Lorenzo Rosasco^{*,†,§} and Tomaso Poggio^{*,¶}

^{*}*Center for Brains, Minds and Machines
McGovern Institute for Brain Research, MIT
43 Vassar St, Cambridge, MA 02139, USA*

[†]*MaLGA, DIBRIS, Università di Genova, Italy*

[‡]*arangam@mit.edu*

[§]*lorenzo.rosasco@unige.it*

[¶]*tp@ai.mit.edu; tp@csail.mit.edu*

Received 17 September 2022

Accepted 13 November 2022

Published 28 December 2022

In this paper, we study kernel ridge-less regression, including the case of interpolating solutions. We prove that maximizing the leave-one-out (CV_{loo}) stability minimizes the expected error. Further, we also prove that the minimum norm solution — to which gradient algorithms are known to converge — is the most stable solution. More precisely, we show that the minimum norm interpolating solution minimizes a bound on CV_{loo} stability, which in turn is controlled by the smallest singular value, hence the condition number, of the empirical kernel matrix. These quantities can be characterized in the asymptotic regime where both the dimension (d) and cardinality (n) of the data go to infinity (with $\frac{n}{d} \rightarrow \gamma$ as $d, n \rightarrow \infty$). Our results suggest that the property of CV_{loo} stability of the learning algorithm with respect to perturbations of the training set may provide a more general framework than the classical theory of Empirical Risk Minimization (ERM). While ERM was developed to deal with the *classical regime* in which the architecture of the learning network is fixed and $n \rightarrow \infty$, the *modern regime* focuses on interpolating regressors and overparameterized models, when both d and n go to infinity. Since the stability framework is known to be equivalent to the classical theory in the classical regime, our results here suggest that it may be interesting to extend it beyond kernel regression to other overparameterized algorithms such as deep networks.

Keywords: Algorithmic stability; kernel regression; minimum norm interpolation; high dimensional statistics; overparameterization.

Mathematics Subject Classification 2020: 68Q32, 68T05

1. Introduction

Statistical learning theory studies the learning properties of machine learning algorithms, and more fundamentally, the conditions under which learning from finite

[¶]Corresponding author.

data is possible. In this context, classical learning theory focuses on the size of the hypothesis space in terms of different complexity measures, such as combinatorial dimensions, covering numbers and Rademacher/Gaussian complexities [6, 26]. Another, more recent, approach is based on defining suitable notions of stability with respect to perturbation of the data [7, 11]. In this view, the continuity of the process that maps data to estimators is crucial, rather than the complexity of the hypothesis space. Different notions of stability can be analyzed, depending on the data perturbation and considered error metric [11]. Interestingly, the stability and complexity approaches to characterizing the *learnability* of problems are not at odds with each other, and have been shown to be equivalent in the classical framework, as shown in [23, 27].

In modern machine learning overparameterized models, with a number of parameters larger than the size of the training data, have now become increasingly common. The ability of these models to generalize is well explained by classical statistical learning theory as long as some form of regularization is used in the training process [8, 28]. However, it was recently shown — first for deep networks [29], and more recently for kernel methods [4, 5] — that learning is possible in the absence of regularization, i.e. when perfectly fitting/interpolating the data. Recent work in statistical learning theory has tried to find theoretical ground for this empirical finding. Since learning using models that interpolate is not exclusive to deep neural networks, we study generalization in the presence of interpolation in the case of kernel methods, with linear models as a special case.

Our Contributions:

- We characterize the generalization properties of possibly interpolating kernel ridge-less regression using a stability approach. While the (uniform) stability properties of regularized kernel methods are well known [7], we study unregularized (“ridgeless”) regression problems.
- We obtain an upper bound on the leave-one-out stability β_{CV} (defined later) of solutions to the kernel least squares problem, and show that this upper bound is minimized by the minimum norm interpolating solution. This also means that among all interpolating solutions, the minimum norm solution has the best test error.^a In particular, the same conclusion is also true for gradient descent and stochastic gradient descent, since these algorithms converge to the minimum norm solution in the setting we consider, see e.g., [25].
- Our stability bounds show that the average stability of the minimum norm solution can be controlled by the minimum eigenvalue of the empirical kernel matrix. It is well known that the numerical stability of the least squares solution is governed by the condition number of the associated kernel matrix which is closely

^aThis holds unless additional information is available, for instance about the data.

related to the minimum eigenvalue (see the discussion of why overparameterization is “good” in [22]). Our results show that the condition number also controls stability (and hence test error) in a statistical sense.

Paper Outline: The rest of the paper is organized as follows. In Sec. [2] we introduce basic ideas in statistical learning and empirical risk minimization, as well as the notation used in the rest of the paper. In Sec. [3] we briefly recall some definitions of stability and their connection to test error. In this section, we also provide a brief discussion about the promise of stability as a framework for the analysis of learning algorithms. In Sec. [4] we present our main results on the stability of kernel least squares. The proof of our theorem is developed in Sec. [6] where we also show that the minimum norm solutions minimize an upper bound on the stability. In Sec. [5] we discuss our results in the context of recent work on high dimensional regression. We support our theoretical results with simulations in Sec. [7] and conclude in Sec. [8].

2. Statistical Learning and Empirical Risk Minimization

We begin by recalling the basic ideas in statistical learning theory. In this setting, X is the input space, Y is the output space, and there is an unknown probability distribution μ on $Z = X \times Y$. In the following, we consider $X = \mathbb{R}^d$ and $Y = \mathbb{R}$. The distribution μ is fixed but unknown, and we are given a training set S consisting of n samples (thus $|S| = n$) drawn i.i.d. from the probability distribution on Z^n , $S = (z_i)_{i=1}^n = (\mathbf{x}_i, y_i)_{i=1}^n$. Intuitively, the goal of supervised learning is to use the training set S to “learn” a function f_S that evaluated at a new value $\mathbf{x}_{\text{new}}$ should predict the associated value of y_{new} , i.e. $y_{\text{new}} \approx f_S(\mathbf{x}_{\text{new}})$.

The loss is a function $V: \mathcal{F} \times Z \rightarrow [0, \infty)$, where $\mathcal{F}$ is the space of measurable functions from X to Y , that measures how well a function performs on a data point. We define a hypothesis space $\mathcal{H} \subseteq \mathcal{F}$ where algorithms search for solutions. With the above notation, the *expected risk* of f is defined as $I[f] = \mathbb{E}_z V(f, z)$ which is the expected loss on a new sample drawn according to the data distribution μ . In this setting, statistical learning can be seen as the problem of finding an approximate minimizer of the expected risk given a training set S . A classical approach to derive an approximate solution is empirical risk minimization (ERM) where we minimize the empirical risk $I_S[f] = \frac{1}{n} \sum_{i=1}^n V(f, z_i)$.

A natural error measure for our ERM solution f_S is the expected excess risk $\mathbb{E}_S[I[f_S] - \min_{f \in \mathcal{H}} I[f]]$. Another common error measure is the expected generalization error/gap given by $\mathbb{E}_S[I[f_S] - I_S[f_S]]$. These two error measures are closely related since the expected excess risk is easily bounded by the expected generalization error (see Lemma [3.1]).

2.1. Kernel least squares and minimum norm solution

The focus in this paper is kernel least squares. We assume the loss function V is the square loss, that is, $V(f, z) = (y - f(\mathbf{x}))^2$. The hypothesis space is assumed to be a

reproducing kernel Hilbert space, defined by a positive definite kernel $K : X \times X \rightarrow \mathbb{R}$ with $\Phi : X \rightarrow \mathcal{H}$ an associated feature map, such that $K(\mathbf{x}, \mathbf{x}') = \langle \Phi(\mathbf{x}), \Phi(\mathbf{x}') \rangle_{\mathcal{H}}$ for all $\mathbf{x}, \mathbf{x}' \in X$, where $\langle \cdot, \cdot \rangle_{\mathcal{H}}$ is the inner product in $\mathcal{H}$. In this setting, functions are linearly parameterized, that is there exists $\mathbf{w} \in \mathcal{H}$ such that $f(\mathbf{x}) = \langle \mathbf{w}, \Phi(\mathbf{x}) \rangle_{\mathcal{H}}$ for all $\mathbf{x} \in X$.

The ERM problem typically has multiple solutions, one of which is the minimum norm solution:

$$f_S^\dagger = \arg \min_{f \in \mathcal{M}} \|f\|_{\mathcal{H}}, \quad \mathcal{M} = \arg \min_{f \in \mathcal{H}} \frac{1}{n} \sum_{i=1}^n (f(\mathbf{x}_i) - y_i)^2, \quad (2.1)$$

where $\|\cdot\|_{\mathcal{H}}$ is the norm in $\mathcal{H}$. The minimum norm solution is unique and satisfies a representer theorem, for all $\mathbf{x} \in X$:

$$f_S^\dagger(\mathbf{x}) = \sum_{i=1}^n K(\mathbf{x}, \mathbf{x}_i) \mathbf{c}_{S,i}, \quad \mathbf{c}_S = \mathbf{K}^\dagger \mathbf{y}, \quad (2.2)$$

where $\mathbf{c}_S = (\mathbf{c}_{S,1}, \dots, \mathbf{c}_{S,n})$, $\mathbf{y} = (y_1 \dots y_n) \in \mathbb{R}^n$, $\mathbf{K}$ is the n by n matrix with entries $\mathbf{K}_{ij} = K(\mathbf{x}_i, \mathbf{x}_j)$, $i, j = 1, \dots, n$ and $\mathbf{K}^\dagger$ is the Moore–Penrose pseudoinverse of $\mathbf{K}$. Since the input points are typically distinct, it is possible to show that for many kernels one can replace $\mathbf{K}^\dagger$ by $\mathbf{K}^{-1}$ (see Remark 2.2). Note that invertibility is necessary and sufficient for interpolation: if $\mathbf{K}$ is invertible, $f_S^\dagger(\mathbf{x}_i) = y_i$ for all $i = 1, \dots, n$, in which case the training error in (2.1) is zero.

An alternative to using the explicit representation of the kernel matrix $\mathbf{K}$ is to represent linear functions in the RKHS as $f(\mathbf{x}) = \langle \mathbf{w}, \Phi(\mathbf{x}) \rangle$ for $\mathbf{w} \in \mathcal{H}$. If we collect the RKHS features of the data $\Phi(\mathbf{x}_i)$ into the rows of a linear operator $\mathbf{X} : \mathcal{H} \rightarrow \mathbb{R}^n$, then we can write the Kernel least square problem as $\min_{\mathbf{w} \in \mathcal{H}} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|_2^2$. All interpolating solutions to this problem are of the form $\hat{\mathbf{w}}_S = \mathbf{X}^\dagger \mathbf{y} + (\mathbf{I}_{\mathcal{H}} - \mathbf{X}^\dagger \mathbf{X})\mathbf{v}$ for any $\mathbf{v} \in \mathcal{H}$. The relationship between the kernel matrix $\mathbf{K}$ and the operator $\mathbf{X}$ is $\mathbf{K} = \mathbf{X}\mathbf{X}^\top$.

Remark 2.1 (Pseudoinverse for Underdetermined Linear Systems). A simple, relevant example is the linear kernel where $f(\mathbf{x}) = \mathbf{w}^\top \mathbf{x}$, $\mathcal{H} = \mathbb{R}^d$ and Φ is the identity map. If the rank of $\mathbf{X} \in \mathbb{R}^{n \times d}$ is n , then any interpolating solution $\mathbf{w}_S$ satisfies $\mathbf{w}_S^\top \mathbf{x}_i = y_i$ for all $i = 1, \dots, n$, and the minimum norm solution, also called Moore–Penrose solution, is given by $\mathbf{w}_S^\dagger = \mathbf{X}^\dagger \mathbf{y}$ where the pseudoinverse $\mathbf{X}^\dagger$ takes the form $\mathbf{X}^\dagger = (\mathbf{X}^\top \mathbf{X})^\dagger \mathbf{X}^\top$.

Remark 2.2 (Invertibility of Translation Invariant Kernels). Translation invariant kernels are a family of kernel functions given by $K(\mathbf{x}_1, \mathbf{x}_2) = k(\mathbf{x}_1 - \mathbf{x}_2)$ where k is an even function on $\mathbb{R}^d$. Translation invariant kernels are Mercer kernels (positive semidefinite) if the Fourier transform of $k(\cdot)$ is non-negative. For Radial Basis Function kernels ($K(\mathbf{x}_1, \mathbf{x}_2) = k(\|\mathbf{x}_1 - \mathbf{x}_2\|)$) we have the additional property due to [18, Theorem 2.3] that for distinct points $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in \mathbb{R}^d$ the kernel matrix $\mathbf{K}$ is non-singular and thus invertible.

The above discussion is directly related to regularization approaches.

Remark 2.3 (Stability and Tikhonov Regularization). Tikhonov regularization is used to prevent potential unstable behaviors. In the above setting, it corresponds to replacing Problem (2.1) by $\min_{f \in \mathcal{H}} \frac{1}{n} \sum_{i=1}^n (f(\mathbf{x}_i) - y_i)^2 + \lambda \|f\|_{\mathcal{H}}^2$ where the corresponding unique solution is given by $f_S^\lambda(\mathbf{x}) = \sum_{i=1}^n K(\mathbf{x}, \mathbf{x}_i) \mathbf{c}_i$, $\mathbf{c} = (\mathbf{K} + \lambda n \mathbf{I}_n)^{-1} \mathbf{y}$. In contrast to ERM solutions, the above approach prevents interpolation. The properties of the corresponding estimator are well known. In this paper, we complement these results focusing on the case $\lambda \rightarrow 0$.

Finally, we end our introductory remarks by recalling the connection between minimum norm and the gradient descent.

Remark 2.4 (Minimum Norm and Gradient Descent). In our setting, it is well known that both batch and stochastic gradient (SGD) iterations converge to the minimum norm solution when multiple solutions exist, see e.g., [25]. Thus, a study of the properties of the minimum norm solution explains the properties of the solution to which SGD converges. In particular, when ERM has multiple interpolating solutions, gradient descent converges to a solution minimizing a bound on stability, as we show next.

3. Error Bounds via Stability

In this section, we present the definition of stability that we will be using in the paper, and discuss how stability may be a unifying framework for explaining learning in both the classical and modern regimes.

We first recall some basic results relating the learning and stability properties of Empirical Risk Minimization (ERM). Throughout the paper, we assume that ERM achieves a minimum, albeit the extension to almost minimizers is possible [19] and important for exponential-type loss functions [21]. We do not require that a minimum exists for the expected risk. Since we will be considering leave-one-out stability in this section, we look at solutions to ERM over the complete training set $S = \{z_1, z_2, \dots, z_n\}$ and the leave one out training set $S_{-i} = \{z_1, z_2, \dots, z_{i-1}, z_{i+1}, \dots, z_n\}$.

The excess risk of ERM can be easily related to its stability properties. Here, we follow the definition laid out in [19] and say that an algorithm is Cross-Validation leave-one-out (CV_{loo}) stable in expectation, if there exists $\beta_{\text{CV}} > 0$ such that for all $i = 1, \dots, n$,

$$\mathbb{E}_S[V(f_{S_{-i}}, z_i) - V(f_S, z_i)] \leq \beta_{\text{CV}}. \quad (3.1)$$

Here $f_S, f_{S_{-i}}$ are the ERM solutions obtained on the full dataset and the leave one out dataset, respectively. This definition is justified by the following result that bounds the excess risk of a learning algorithm by its average CV_{loo} stability [19, 27].

Lemma 3.1 (Excess Risk and CV_{loo} Stability). For all $i = 1, \dots, n$,

$$\mathbb{E}_S \left[I[f_{S_{-i}}] - \inf_{f \in \mathcal{H}} I[f] \right] \leq \mathbb{E}_S[V(f_{S_{-i}}, z_i) - V(f_S, z_i)]. \quad (3.2)$$

Remark 3.1 (Connection to Other Notions of Stability). Uniform stability, introduced by [7], corresponds in our notation to the assumption that there exists $\beta_u > 0$ such that for all $i = 1, \dots, n$, $\sup_{z \in Z} |V(f_{S_{-i}}, z) - V(f_S, z)| \leq \beta_u$. Clearly, this is a strong notion implying most other definitions of stability. We note that there are a number of different notions of stability. We refer the interested reader to [11, 19].

Lemma 3.1 is known and we recall the proof in Appendix A.2 for completeness. In Appendix A we also discuss other definitions of stability and their connections to concepts in statistical learning theory like generalization and learnability.

3.1. The stability framework: A unifying principle for the classical and modern regimes

A milestone in classical learning theory was to formally show that appropriately restricting the hypothesis space — that is the space of functions represented by the learning machine — ensures consistency (and generalization) of ERM. The classical theory assumes that the hypothesis space is fixed while the number of training data n increases to infinity. Its basic results thus characterize the “classical” regime of $n > d$, where d is the number of parameters to be learned. The classical theory, however, cannot deal with what we call the “modern” regime, in which the network remains overparameterized ($n < d$) when n grows. In this case, the hypothesis space is not fixed: d increases as n increases. Different approaches that do not rely on the hypothesis space were developed already twenty years ago, motivated by learning algorithms that are not ERM, such as k -Nearest Neighbor. While trying to develop a theory that can deal with the classical *and* the modern regime, it seems natural to abandon the idea of the hypothesis space as the object of interest and focus instead on properties of supervised learning algorithms, which are maps from data sets to hypothesis functions. One can ask: *what property must the learning map L have for good generalization error?* The answer for a fixed hypothesis space is that CV_{loo} stability is necessary and sufficient for generalization and consistency of ERM.^b

Building upon this observation, we conjecture that CV_{loo} stability may be used to develop a unifying theory encompassing both the classical and the modern regime for ERM. In the classical regime generalization can take place provided $\beta_{\text{CV}} \rightarrow 0$ when $n \rightarrow \infty$; for ERM consistency follows from generalization. In the modern interpolatory regime, the generalization gap given by $\mathbb{E}_S[I[f_S] - I_S[f_S]]$ does not necessarily decrease to 0 as $n \rightarrow \infty$ since $I_S[f_S] = 0$, while we can have $I[f_S] > 0$ in general. However, for interpolating regressors, β_{CV} becomes a bound on the expected error. Thus, the key claim of this unified approach is that *minimizing β_{CV} minimizes the expected generalization gap* and in particular minimizes the

^bLOO stability (see [23]) together with CV_{loo} stability of the algorithm, both going to zero for $n \rightarrow \infty$ is sufficient for generalization for any supervised algorithm, including k -nearest neighbor and kernel machines.

expected error in the modern regime. While this is satisfying conceptually, it is also important to spell out the implications of minimizing β_{CV} for ERM.^c A natural answer is that minimizing β_{CV} in ERM may be equivalent to selecting the minimum norm solution. For the case of kernel regressors, we show next that the minimum norm ERM interpolator indeed minimizes CV_{loo} stability. It should be emphasized that this is an upper bound and we cannot expect the minimum norm solution to always yield the minimum expected error. In particular, better solution can be found when prior information is available (see [20]). In addition, it remains an open question whether similar results hold for classifiers such as deep networks.^d

4. CV_{loo} Stability of Kernel Least Squares

In this section, we analyze the expected CV_{loo} stability of solutions to the kernel least squares problem, and obtain a corresponding upper bound on their stability. We show that the upper bound on the expected CV_{loo} stability is governed by the norm of the solutions in the case of interpolating solutions, and hence is the smallest for the minimum norm interpolating solution (2.1).

As outlined in Sec. 2.1 we consider a kernel least squares problem on a dataset $S = \{(\mathbf{x}_i, y_i)\}_{i=1}^n \subseteq (\mathbb{R}^d \times \mathbb{R})^n$. We use the linear parameterization in the RKHS $f(\mathbf{x}) = \langle \mathbf{w}, \Phi(\mathbf{x}) \rangle_{\mathcal{H}}$, and collect the RKHS features of the data $\Phi(\mathbf{x}_i)$ into the rows of a linear operator $\mathbf{X}: \mathcal{H} \rightarrow \mathbb{R}^n$. We recall that all interpolating solutions are of the form $\hat{\mathbf{w}}_S = \mathbf{X}^\dagger \mathbf{y} + (\mathbf{I}_{\mathcal{H}} - \mathbf{X}^\dagger \mathbf{X})\mathbf{v}$ for some $\mathbf{v} \in \mathcal{H}$. Since we are interested in CV_{loo} stability, we consider the same kernel least squares problem on the leave one out dataset S_{-i} . We replace the i^{th} row of $\mathbf{X}$ with $\mathbf{0}_{\mathcal{H}}$ to obtain the corresponding data operator $\mathbf{X}_{-i}: \mathcal{H} \rightarrow \mathbb{R}^n$ for the leave one out dataset. All solutions on the leave one out dataset S_{-i} can be written as $\hat{\mathbf{w}}_{S_{-i}} = (\mathbf{X}_{-i})^\dagger \mathbf{y}_{-i} + (\mathbf{I}_{\mathcal{H}} - \mathbf{X}_{-i}^\dagger \mathbf{X}_{-i})\mathbf{v}_{-i}$ for some $\mathbf{v}_{-i} \in \mathcal{H}$. We note that when $\mathbf{v} = \mathbf{0}_{\mathcal{H}}$ and $\mathbf{v}_{-i} = \mathbf{0}_{\mathcal{H}}$, we obtain the minimum norm interpolating solutions on the datasets S and S_{-i} .

Theorem 4.1 (Main Theorem). *Consider the kernel least squares problem where the inputs $\mathbf{x} \in \mathcal{H}$ and the outputs y are bounded, that is there exist $\kappa, M > 0$ such that*

$$\|\mathbf{x}\|_{\mathcal{H}}^2 \leq \kappa, \quad |y| \leq M, \quad (4.1)$$

almost surely. Then for any interpolating solutions $\hat{f}_{S_{-i}}, \hat{f}_S$,

$$\mathbb{E}_S[V(\hat{f}_{S_{-i}}, z_i) - V(\hat{f}_S, z_i)] \leq C_1 \mathbb{E}_S[\beta_{CV}] + C_2 \mathbb{E}_S[\beta_{CV}^2], \quad (4.2)$$

^cAn argument may be made that while overparameterization makes sense for large but finite amounts of data, it should disappear for realistic learning machines as $n \rightarrow \infty$. If this does not happen the empirical loss will never converge to the expected loss, which seems a natural requirement, especially in a quasi-online setting.

^dThe results of course hold for deep RELU networks in the NTK regime, since they are then equivalent to kernel machines. Our results provide context for NTK analyzes similar to that presented in [1].

where $\beta_{CV} = \|\mathbf{X}^\dagger\|_{\text{op}}\|\mathbf{y}\| + 2\|\mathbf{v} - \mathbf{v}_{-i}\| + \|\mathbf{v}_{-i}\|$ and C_1, C_2 are absolute constants that do not depend on either d or n . This bound is minimized when $\mathbf{v} = \mathbf{v}_{-i} = \mathbf{0}_{\mathcal{H}}$, which corresponds to the minimum norm interpolating solutions $f_S^\dagger, f_{S_{-i}}^\dagger$. For the minimum norm solutions we have $\beta_{CV}^{\min} = \|\mathbf{X}^\dagger\|_{\text{op}}\|\mathbf{y}\|$.

In the above theorem $\|\mathbf{X}^\dagger\|_{\text{op}}$ refers to the operator norm of the pseudoinverse of the data operator $\mathbf{X}$, $\|\mathbf{y}\|$ refers to the (Euclidean) norm of $\mathbf{y} \in \mathbb{R}^n$.

We can combine the above result with Lemma 3.1 to obtain the following bound on excess risk for minimum norm interpolating solutions to the kernel least squares problem.

Corollary 4.1. *The excess risk of the minimum norm interpolating kernel least squares solution can be bounded as*

$$\mathbb{E}_S \left[I[f_{S_{-i}}^\dagger] - \inf_{f \in \mathcal{H}} I[f] \right] \leq C_1 \mathbb{E}_S [\|\mathbf{X}^\dagger\|_{\text{op}}\|\mathbf{y}\|] + C_2 \mathbb{E}_S [\|\mathbf{X}^\dagger\|_{\text{op}}^2 \|\mathbf{y}\|^2].$$

We provide the proof of Theorem 4.1 in Sec. 6. In the next section, we first offer some discussion of our results on stability, and put our results in the context of other recent results on interpolation in linear and kernel least squares problems.

5. Discussion and Related Work

In the previous section, we obtained bounds on the CV_{loo} stability of kernel least squares solutions and in particular of *interpolating* solutions. We established a bound on average stability for kernel least squares solutions, and show that this bound is minimized when the minimum norm ERM solution is selected. One of our key findings is the relationship between minimizing the norm of the ERM solution and minimizing a bound on stability. In this section, we discuss our bound under different regimes of the sample size n and the dimensionality of the data d .

For the kernel least squares problem, interpolation occurs under mild conditions for different kernels. For instance, if the input data are all distinct, the inverse of the kernel matrix exists and for positive definite radial kernels interpolation is expected. For other kernels, such as the linear kernel, $d \geq n$ is needed for interpolation.

Asymptotic Analysis: While our bounds hold for any finite d and n , it is worth understanding how they evolve under different regimes of n and d . For $n \rightarrow \infty$ (and d fixed), the smallest singular value of the kernel matrix $\mathbf{K} = [K(\mathbf{x}_i, \mathbf{x}_j)]$ typically decreases with n . This means our bounds diverge and stability is lost. The classical approach here is to use regularized ERM (see Remark 2.3) corresponding to

$$\min_{f \in \mathcal{H}} \frac{1}{n} \sum_{i=1}^n (f(\mathbf{x}_i) - y_i)^2 + \lambda \|f\|_K^2, \quad (5.1)$$

which gives the following set of equations for $\mathbf{c}$ (with $\lambda \geq 0$):

$$(\mathbf{K} + n\lambda\mathbf{I})\mathbf{c} = \mathbf{y}. \quad (5.2)$$

Regularized ERM has a strong stability guarantee with a uniform stability bound (defined in Appendix A.3) $\beta = O(\frac{1}{\lambda n})$ which turns out to be inversely proportional to the regularization parameter λ and the sample size n [7]. Here the limit $n \rightarrow \infty$ implies asymptotic convergence to a zero stability gap.

In a setting that is common in statistics, we can also consider how our bounds evolve as both the n and d go to infinity, but the ratio $\frac{n}{d} \rightarrow \gamma$ remains finite as $n, d \rightarrow \infty$. In this setting, it is possible to use results from random matrix theory [15] to sketch the asymptotic limits of our bound for linear kernels under distributional assumptions on the data. Since $\|\mathbf{X}^\dagger\|_{\text{op}} = \frac{1}{\sigma_{\min}(\mathbf{X})}$, we can compute the asymptotic limit of our bound as $\|\mathbf{X}^\dagger\|_{\text{op}}^2 \|\mathbf{y}\|^2 = \frac{n}{d(1-\sqrt{\gamma})^2} = \frac{\gamma}{(1-\sqrt{\gamma})^2}$. Notice that the bound does not go to zero for $n \rightarrow \infty$ because in general the expected error cannot vanish (unless the classification labels are deterministic). Since the kernel matrix $\mathbf{K}$ is related to the data operator $\mathbf{X}$ as $\mathbf{K} = \mathbf{X}^\top \mathbf{X}$, we have that $\sigma_{\min}(\mathbf{K}) = \sigma_{\min}(\mathbf{X})^2$, and our bound can be written in terms of the minimum singular value of the kernel matrix rather than data operator.

Interestingly, properties analogous to the Marchenko–Pastur limit hold for more general kernels. Consider random matrices whose entries are $K(\mathbf{x}_i^\top \mathbf{x}_j)$ with i.i.d. vectors $\mathbf{x}_i$ in $\mathbb{R}^d$ with mean zero and unit variance. Assuming that the distribution of $\mathbf{x}_i$'s is sufficiently nice and f is sufficiently smooth, [9] showed that in the Marchenko–Pastur limit, the spectral distributions of kernel dot-product matrices $\mathbf{K}_{ij} = f(\frac{1}{d}\mathbf{x}_i^\top \mathbf{x}_j)$ behave as if f is linear. In fact, El Karoui showed that under mild conditions, the kernel matrix is asymptotically equivalent to a linear combination of the linear kernel matrix, the all-1's matrix, and the identity, and hence the limiting spectrum is Marcenko–Pastur.^e However, we note that our results *do not* predict a double descent curve for kernels that are not linear dot product kernels [22]. We discuss this observation in more detail in Sec. 7.

Related Work: Recently, there has been a surge of interest in studying linear and kernel least squares models, since classical results focus on situations where constraints or penalties that prevent interpolation are added to the empirical risk. For example, high dimensional linear regression is considered in [10, 16] from the perspective of asymptotic statistics. A non-asymptotic approach is considered instead in [3, 12, 13, 24]. In particular, the results in [3] are the first to obtain convergence when the number of dimensions/parameters is fixed.

While these papers study upper and lower bounds on the risk of interpolating solutions to the linear and kernel least squares problem, ours are the first to be derived using stability arguments. While it might be possible to obtain tighter excess risk bounds through careful analysis of the minimum norm interpolant, our simple approach helps us establish a link between stability in a statistical and numerical sense. Of course, our result is in terms of an upper bound and since

^eRemark 5.1 of [13] observes that since the data is usually centered ($\sum_{i=1}^n \mathbf{x}_i = \mathbf{0}$), the spectrum of the kernel matrix is close to the spectrum of the linear kernel.

lower bounds do not yet exist and seem difficult to obtain, it is reasonable to be skeptical of its quantitative values. More relevant in our opinion is the qualitative statement that minimizing the norm of an interpolating solution has the effect of making its stability gap smaller and thus of minimizing its expected error. We also see this reflected in numerical simulations in Sec. 7. Concurrent to our work, [14] study the classification problem using kernels and obtain a mistake bound for the minimum norm interpolating classifier. However, they do not make the connection to CV_{loo} stability.

Finally, we can compare our results with observations made in [22] on the condition number of random kernel matrices. The condition number of the empirical kernel matrix is known to control the numerical stability of the solution to a kernel least squares problem. Our results show that the statistical stability is also controlled by the minimum singular value of the kernel matrix (which is closely related to the condition number), providing a natural link between numerical and statistical stability.

6. Proof of Theorem 4.1

6.1. Key lemmas

In order to prove Theorem 4.1 we make use of the following lemmas to bound the CV_{loo} stability using the norms and the difference of the solutions.

Lemma 6.1. *Under assumption (4.1), for all $i = 1, \dots, n$, it holds that*

$$\mathbb{E}_S[V(\hat{f}_{S_{-i}}, z_i) - V(\hat{f}_S, z_i)] \leq \mathbb{E}_S[(2M + \kappa(\|\hat{f}_S\|_{\mathcal{H}} + \|\hat{f}_{S_{-i}}\|_{\mathcal{H}})) \times \kappa\|\hat{f}_S - \hat{f}_{S_{-i}}\|_{\mathcal{H}}].$$

Proof. We begin, recalling that the square loss is locally Lipschitz, that is for all $y, a, a' \in \mathbb{R}$, with

$$|(y - a)^2 - (y - a')^2| \leq (2|y| + |a| + |a'|)|a - a'|.$$

If we apply this result to f, f' in a RKHS $\mathcal{H}$,

$$|(y - f(\mathbf{x}))^2 - (y - f'(\mathbf{x}))^2| \leq \kappa(2M + \kappa(\|f\|_{\mathcal{H}} + \|f'\|_{\mathcal{H}}))\|f - f'\|_{\mathcal{H}}.$$

Using the basic properties of a RKHS that for all $f \in \mathcal{H}$,

$$|f(\mathbf{x})| \leq \|f\|_{\infty} \leq \kappa\|f\|_{\mathcal{H}}. \quad (6.1)$$

In particular, we can plug $\hat{f}_{S_{-i}}$ and $\hat{f}_S$ into the above inequality, and the almost positivity of ERM [19] will allow us to drop the absolute value on the left-hand side. Finally, the desired result follows by taking the expectation over S . $\square$

Now that we have bounded the CV_{loo} stability using the norms and the difference of the solutions, we can find a bound on the difference between the solutions to the kernel least squares problem. This is our main stability estimate.

Lemma 6.2. *Let $\hat{f}_S, \hat{f}_{S_{-i}}$ be any interpolating kernel least squares solutions on the full and leave one out datasets (as defined at the top of this section), then*

$\|\hat{f}_S - \hat{f}_{S_{-i}}\|_{\mathcal{H}} \leq \|\mathbf{X}^\dagger\|_{\text{op}} \|\mathbf{y}\| + 2\|\mathbf{v} - \mathbf{v}_{-i}\| + \|\mathbf{v}_{-i}\|$. This bound is minimized when $\mathbf{v} = \mathbf{v}_{-i} = \mathbf{0}_{\mathcal{H}}$, which corresponds to the minimum norm interpolating solutions $f_S^\dagger, f_{S_{-i}}^\dagger$.

Also,

$$\|f_S^\dagger - f_{S_{-i}}^\dagger\| \leq \|\mathbf{X}^\dagger\|_{\text{op}} \|\mathbf{y}\|. \quad (6.2)$$

Remark 6.1 (Zero Training Loss). In Lemma 6.1 we use the locally Lipschitz property of the squared loss function to bound the leave one out stability in terms of the difference between the norms of the solutions. Under interpolating conditions, if we set the term $V(\hat{f}_S, z_i) = 0$, the leave one out stability reduces to $\mathbb{E}_S[V(\hat{f}_{S_{-i}}, z_i) - V(\hat{f}_S, z_i)] = \mathbb{E}_S[V(\hat{f}_{S_{-i}}, z_i)] = \mathbb{E}_S[(\hat{f}_{S_{-i}}(\mathbf{x}_i) - y_i)^2] = \mathbb{E}_S[(\hat{f}_{S_{-i}}(\mathbf{x}_i) - \hat{f}_S(\mathbf{x}_i))^2] = \mathbb{E}_S[\langle \hat{f}_{S_{-i}}(\cdot) - \hat{f}_S(\cdot), K_{\mathbf{x}_i}(\cdot) \rangle^2] \leq \mathbb{E}_S[\|\hat{f}_S - \hat{f}_{S_{-i}}\|_{\mathcal{H}}^2 \times \kappa^2]$. We can plug in the bound from Lemma 6.2 to obtain similar qualitative and quantitative (up to constant factors) results as in Theorem 4.1

Since the minimum norm interpolating solutions minimize both $\|\hat{f}_S\|_{\mathcal{H}} + \|\hat{f}_{S_{-i}}\|_{\mathcal{H}}$ and $\|\hat{f}_S - \hat{f}_{S_{-i}}\|_{\mathcal{H}}$ (from Lemmas 6.1 and 6.2), we can put them together to prove Theorem 4.1. In the following section, we provide the proof of Lemma 6.2

6.2. Proof of Lemma 6.2

We have n samples in the training set for a kernel least squares problem, $\{(\mathbf{x}_i, y_i)\}_{i=1}^n$. We consider the linear operator $\mathbf{X} = [\Phi(\mathbf{x}_1)^\top; \Phi(\mathbf{x}_2)^\top; \dots; \Phi(\mathbf{x}_n)^\top]$ from $\mathcal{H}$ to $\mathbb{R}^n$ and vector of labels $\mathbf{y} = [y_1 y_2 \dots y_n]^\top \in \mathbb{R}^n$. For any $\mathbf{w} \in \mathcal{H}$, the operator evaluates to $\mathbf{X}\mathbf{w} \in \mathbb{R}^n$, the i th entry of which is given by $\langle \mathbf{w}, \Phi(\mathbf{x}_i) \rangle_{\mathcal{H}}$. Then any ERM solution $\mathbf{w}_S$ satisfies the linear equation

$$\mathbf{X}\hat{\mathbf{w}}_S = \mathbf{y}. \quad (6.3)$$

Any solution can be written as

$$\hat{\mathbf{w}}_S = \mathbf{X}^\dagger \mathbf{y} + (\mathbf{I}_{\mathcal{H}} - \mathbf{X}^\dagger \mathbf{X}) \mathbf{v}. \quad (6.4)$$

If we consider the leave one out training set S_{-i} we can find the minimum norm ERM solution for $\mathbf{X}_{-i} = [\Phi(\mathbf{x}_1)^\top; \dots; \mathbf{0}_{\mathcal{H}}^\top; \dots; \Phi(\mathbf{x}_n)^\top]$ and $\mathbf{y}_{-i} = [y_1 \dots 0 \dots y_n]^\top$ as

$$\hat{\mathbf{w}}_{S_{-i}} = (\mathbf{X}_{-i})^\dagger \mathbf{y}_{-i} + (\mathbf{I}_{\mathcal{H}} - \mathbf{X}_{-i}^\dagger \mathbf{X}_{-i}) \mathbf{v}_{-i}. \quad (6.5)$$

We can write $\mathbf{X}_{-i}$ as

$$\mathbf{X}_{-i} = \mathbf{X} + \mathbf{a}\mathbf{b}^\top, \quad (6.6)$$

where $\mathbf{b} \in \mathcal{H}$ is a vector representing the additive change to the i th row, i.e. $\mathbf{b} = -\Phi(\mathbf{x}_i)$, and $\mathbf{a} = \mathbf{e}_i \in \mathbb{R}^n$ is the i th element of the canonical basis in $\mathbb{R}^n$ (all the coefficients are zero but the i th which is one). Thus, $\mathbf{a}\mathbf{b}^\top$ is a linear operator from $\mathcal{H}$ to $\mathbb{R}^n$ that maps vectors in $\mathcal{H}$ to scaled versions of $\mathbf{a}$.

We also have $\mathbf{y}_{-i} = \mathbf{y} - y_i \mathbf{a}$. Now per Lemma 6.1 we are interested in bounding the quantity $\|\hat{f}_{S_{-i}} - \hat{f}_S\|_{\mathcal{H}} = \|\hat{\mathbf{w}}_{S_{-i}} - \hat{\mathbf{w}}_S\|_{\mathcal{H}}$. This simplifies to

$$\begin{aligned}
 \|\hat{\mathbf{w}}_{S_{-i}} - \hat{\mathbf{w}}_S\|_{\mathcal{H}} &= \|\mathbf{X}_{-i}^\dagger \mathbf{y}_{-i} - \mathbf{X}^\dagger \mathbf{y} + \mathbf{v}_{-i} - \mathbf{v} + \mathbf{X}^\dagger \mathbf{X} \mathbf{v} - \mathbf{X}_{-i}^\dagger \mathbf{X}_{-i} \mathbf{v}_{-i}\|_{\mathcal{H}} \\
 &= \|(\mathbf{X}_{-i})^\dagger (\mathbf{y} - y_i \mathbf{a}) - \mathbf{X}^\dagger \mathbf{y} + \mathbf{v}_{-i} - \mathbf{v} + \mathbf{X}^\dagger \mathbf{X} \mathbf{v} - \mathbf{X}_{-i}^\dagger \mathbf{X}_{-i} \mathbf{v}_{-i}\|_{\mathcal{H}} \\
 &= \|(\mathbf{X}_{-i}^\dagger - \mathbf{X}^\dagger) \mathbf{y} + y_i \mathbf{X}_{-i}^\dagger \mathbf{a} + \mathbf{v}_{-i} - \mathbf{v} + \mathbf{X}^\dagger \mathbf{X} \mathbf{v} - \mathbf{X}_{-i}^\dagger \mathbf{X}_{-i} \mathbf{v}_{-i}\| \\
 &= \|(\mathbf{X}_{-i}^\dagger - \mathbf{X}^\dagger) \mathbf{y} + \mathbf{v}_{-i} - \mathbf{v} + \mathbf{X}^\dagger \mathbf{X} \mathbf{v} - \mathbf{X}_{-i}^\dagger \mathbf{X}_{-i} \mathbf{v}_{-i}\| \\
 &= \|(\mathbf{X}_{-i}^\dagger - \mathbf{X}^\dagger) \mathbf{y} + (\mathbf{I}_{\mathcal{H}} - \mathbf{X}^\dagger \mathbf{X})(\mathbf{v}_{-i} - \mathbf{v}) \\
 &\quad + (\mathbf{X}^\dagger \mathbf{X} - \mathbf{X}_{-i}^\dagger \mathbf{X}_{-i}) \mathbf{v}_{-i}\|. \tag{6.7}
 \end{aligned}$$

In the above equation, we make use of the fact that $\mathbf{X}_{-i}^\dagger \mathbf{a} = \mathbf{0}_{\mathcal{H}}$. We use an old formula [2, 17] to compute $(\mathbf{X}_{-i})^\dagger$ from $\mathbf{X}^\dagger$. We use the development of pseudo-inverses of perturbed matrices in [17]. Since none of the theorems depend on the finite dimensionality of $\mathcal{H}$, we can use those results for linear operators. We see that $\mathbf{b} = -\Phi(\mathbf{x}_i)$ is a vector in the range of $\mathbf{X}^\top$ and $\mathbf{a}$ is in the range of $\mathbf{X}$ (provided $\mathbf{X}$ has rank n), with $\beta = 1 + \mathbf{b}^\top \mathbf{X}^\dagger \mathbf{a} = 1 - \Phi(\mathbf{x}_i)^\top \mathbf{X}^\dagger \mathbf{a} = 0$. This means we can use Theorem 6 in [17] (equivalent to formula (2.1) in [2]) to obtain the expression for $\mathbf{X}_{-i}^\dagger$,

$$\mathbf{X}_{-i}^\dagger = \mathbf{X}^\dagger - \mathbf{k} \mathbf{k}^\dagger \mathbf{X}^\dagger - \mathbf{X}^\dagger \mathbf{h} \mathbf{h}^\dagger + (\mathbf{k}^\dagger \mathbf{X}^\dagger \mathbf{h}^\dagger) \mathbf{k} \mathbf{h}, \tag{6.8}$$

where $\mathbf{k} = \mathbf{X}^\dagger \mathbf{a}$ and $\mathbf{h} = \mathbf{b}^\top \mathbf{X}^\dagger$ and $\mathbf{u}^\dagger = \frac{\mathbf{u}^\top}{\|\mathbf{u}\|^2}$ for any nonzero vector $\mathbf{u}$,

$$\begin{aligned}
 \mathbf{X}_{-i}^\dagger - \mathbf{X}^\dagger &= (\mathbf{k}^\dagger \mathbf{X}^\dagger \mathbf{h}^\dagger) \mathbf{k} \mathbf{h} - \mathbf{k} \mathbf{k}^\dagger \mathbf{X}^\dagger - \mathbf{X}^\dagger \mathbf{h} \mathbf{h}^\dagger \\
 &= (\mathbf{k}^\dagger \mathbf{X}^\dagger \mathbf{a}) \mathbf{k} \mathbf{a}^\top - \mathbf{k} \mathbf{k}^\dagger \mathbf{X}^\dagger - \mathbf{X}^\dagger \mathbf{a} \mathbf{a}^\top \\
 &= (\mathbf{k}^\dagger \mathbf{k}) \mathbf{k} \mathbf{a}^\top - \mathbf{k} \mathbf{a}^\top - \mathbf{k} \mathbf{k}^\dagger \mathbf{X}^\dagger, \\
 \Rightarrow \|\mathbf{X}_{-i}^\dagger - \mathbf{X}^\dagger\|_{\text{op}} &= \|\mathbf{k} \mathbf{k}^\dagger \mathbf{X}^\dagger\|_{\text{op}} \\
 &\leq \|\mathbf{X}^\dagger\|_{\text{op}}. \tag{6.9}
 \end{aligned}$$

The above set of inequalities follows from the fact that the operator norm of a rank 1 matrix is given by $\|\mathbf{u} \mathbf{v}^\top\|_{\text{op}} = \|\mathbf{u}\| \times \|\mathbf{v}\|$, and by noticing that $\mathbf{k} = -\mathbf{b}$.

Also, from List 2 of [2], we have that $\mathbf{X}_{-i}^\dagger \mathbf{X}_{-i} = \mathbf{X}^\dagger \mathbf{X} - \mathbf{k} \mathbf{k}^\dagger$.

Plugging in these calculations into Eq. (6.7) we get

$$\begin{aligned}
 \|\hat{\mathbf{w}}_{S_{-i}} - \hat{\mathbf{w}}_S\| &= \|(\mathbf{X}_{-i}^\dagger - \mathbf{X}^\dagger) \mathbf{y} + (\mathbf{I}_{\mathcal{H}} - \mathbf{X}^\dagger \mathbf{X})(\mathbf{v}_{-i} - \mathbf{v}) - (\mathbf{X}^\dagger \mathbf{X} - \mathbf{X}_{-i}^\dagger \mathbf{X}_{-i}) \mathbf{v}_{-i}\| \\
 &\leq \|\mathbf{X}^\dagger\|_{\text{op}} \|\mathbf{y}\| + \|\mathbf{I}_{\mathcal{H}} - \mathbf{X}^\dagger \mathbf{X}\|_{\text{op}} \|\mathbf{v} - \mathbf{v}_{-i}\| + \|\mathbf{k} \mathbf{k}^\dagger\|_{\text{op}} \|\mathbf{v}_{-i}\| \\
 &\leq \|\mathbf{X}^\dagger\|_{\text{op}} \|\mathbf{y}\| + 2\|\mathbf{v} - \mathbf{v}_{-i}\| + \|\mathbf{v}_{-i}\|. \tag{6.10}
 \end{aligned}$$

We see that the right-hand side is minimized when $\mathbf{v} = \mathbf{v}_{-i} = \mathbf{0}_{\mathcal{H}}$. This concludes the proof of Lemma 6.2

Remark 6.2 (Stability of the Minimum Norm Solution). We can perform a more careful analysis of the stability of the minimum norm solution by putting together Eqs. (6.7) and (6.9), with $\mathbf{v} = \mathbf{v}_{-i} = \mathbf{0}_{\mathcal{H}}$ to obtain the following bound:

$$\|\mathbf{w}_{S_{-i}}^{\dagger} - \mathbf{w}_S^{\dagger}\| = \|\mathbf{k}\mathbf{k}^{\dagger}\mathbf{X}^{\dagger}\mathbf{y}\| \leq \|\mathbf{w}_S^{\dagger}\|. \quad (6.11)$$

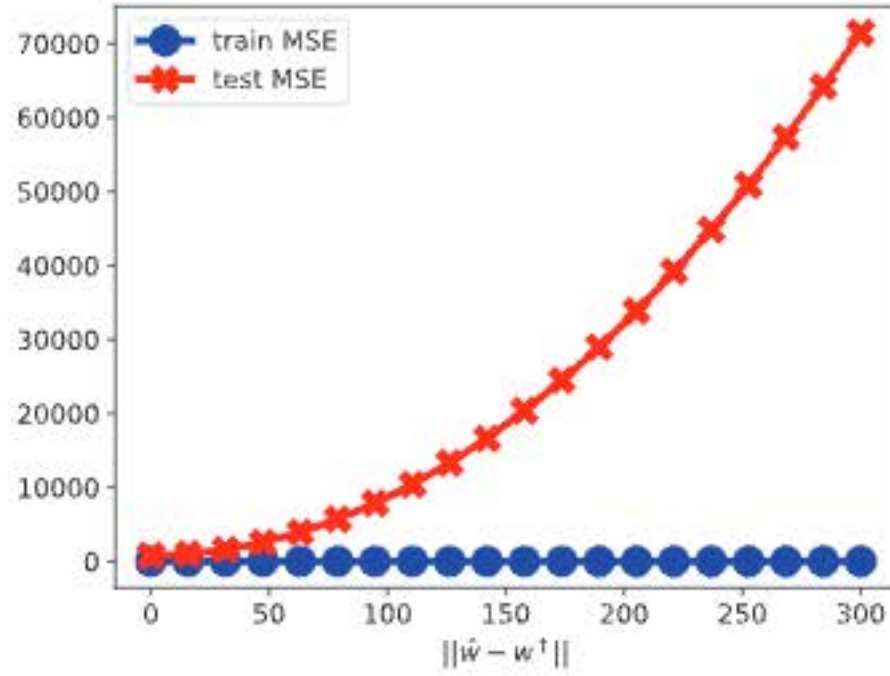
Putting this together with Lemmas 6.1 and 3.1 we can see that the CV_{100} stability — and hence excess risk of the minimum norm solution to the kernel least squares problem — is bounded by the norm of the solution.

7. Simulations

In this section, we perform experiments to provide empirical evidence for our theoretical results. We first verify that the minimum norm interpolating solution maximizes stability among all interpolating solutions. Subsequently, we show that the test error inversely correlates with the minimum singular value of the kernel matrix, and we finally verify that the norm of the minimum norm interpolating solution governs its test error. We will now see those experiments one by one.

Stability and Norm of the Solution: In order to illustrate that the minimum norm interpolating solution is the best performing interpolating solution we run a simple experiment on a linear regression problem. We synthetically generate data from a linear model $\mathbf{y} = \mathbf{X}\mathbf{w}$, where $\mathbf{X} \in \mathbb{R}^{n \times d}$ is i.i.d $\mathcal{N}(0, 1)$. The dimension of the data is $d = 1000$ and there are $n = 200$ samples in the training dataset. We use a held out test dataset of 100 samples to measure the generalization performance. We compute interpolating solutions as described at the beginning of this section, $\hat{\mathbf{w}} = \mathbf{X}^{\dagger}\mathbf{y} + (\mathbf{I} - \mathbf{X}^{\dagger}\mathbf{X})\mathbf{v}$, using $\mathbf{v}$'s of different norms to compare the test error and CV_{100} stability. The results (averaged over 100 trials) are shown in Fig. 1. In Fig. 1(a), we can see that the training loss is 0 for all interpolating solutions, but the test MSE increases as $\|\mathbf{v}\|$ increases, with the minimum norm solution $\mathbf{w}^{\dagger} = \mathbf{X}^{\dagger}\mathbf{y}$ having the best performance. We also observe in Fig. 1(b) that the CV_{100} stability, computed using the expression in (3.1), also follows a similar trend. From both plots in Fig. 1 we can see that the minimum norm interpolating solution has the best stability as well as the best test error, as suggested by Theorem 4.1

Test Error and $\frac{1}{\sigma_{\min}(\mathbf{K})}$: Our results also indicate that the bound on the CV_{100} stability (and hence the test error) of the minimum norm interpolating solution depends on the norm of the (pseudo) inverse of the empirical kernel matrix. In order to verify this using simulations, we consider a regression problem in which we learned the function $f(\mathbf{x}) = \exp(-2\|\mathbf{x}\|)$ using kernel least squares. We generated samples $\mathbf{x}_i \in \mathbb{R}^{20}$ and learned f using interpolating kernel “ridgeless” regression with a polynomial kernel of degree 2 as well as a radial basis function (RBF) kernel. The training dataset was generated by sampling $\mathbf{X} \in \mathbb{R}^{n \times d}$ ($d = 20, n = 200$) i.i.d. from $\mathcal{N}(0, 1)$, and a held out test dataset of 100 samples was also generated in a similar fashion. The results of this simulation can be seen in Fig. 2. In order to



(a) Train and Test MSE

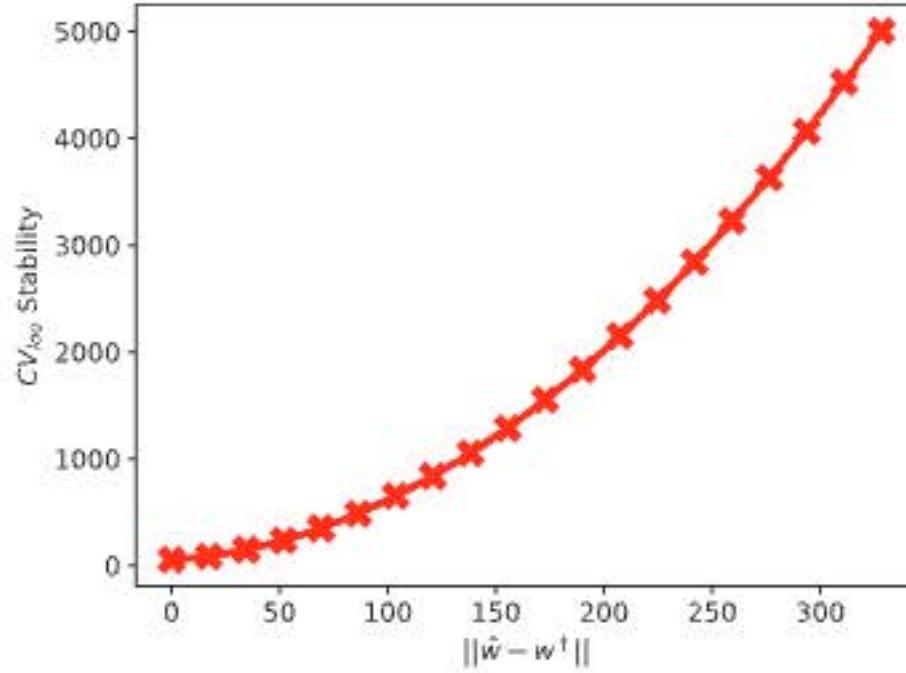
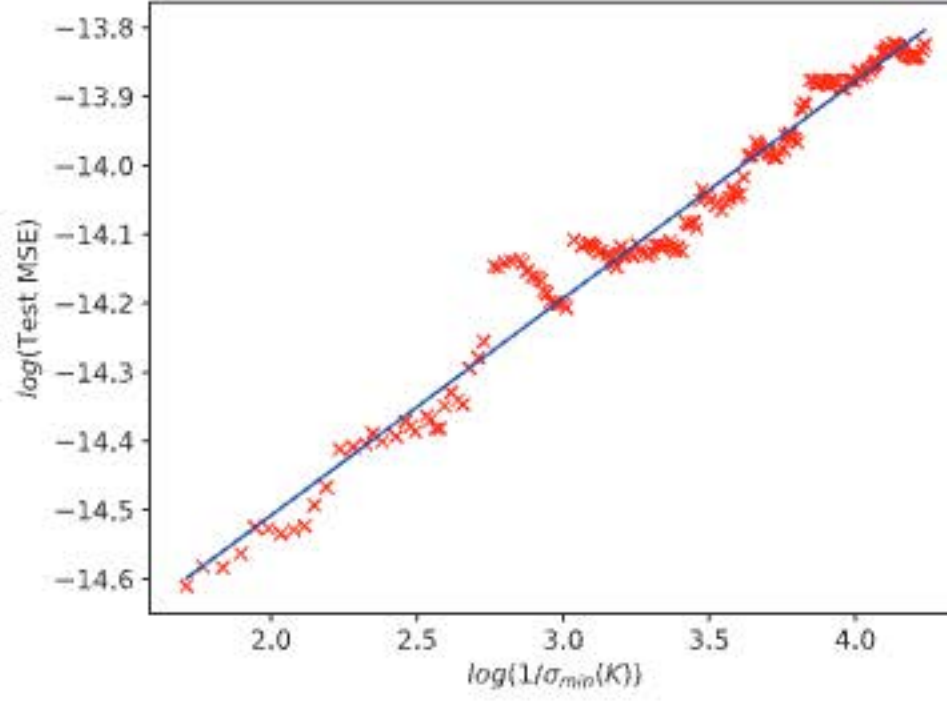
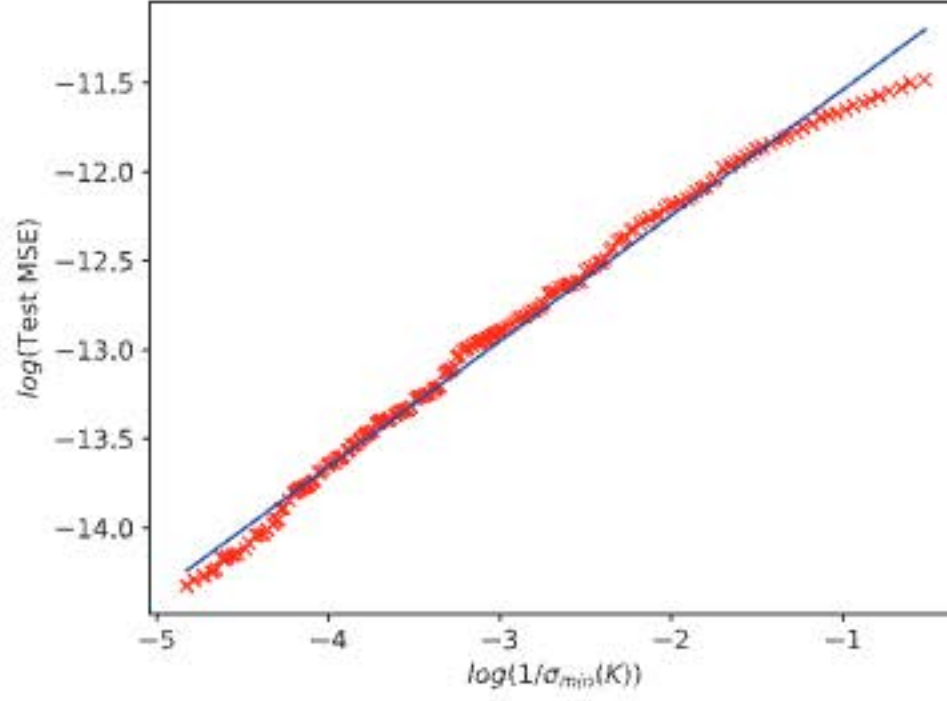
(b) CV_{100} Stability

Fig. 1. Plot of (a) Train and test mean squared error and (b) CV_{100} stability versus distance between an interpolating solution $\hat{\mathbf{w}}$ and the minimum norm interpolating solution $\mathbf{w}^\dagger$ of a linear regression problem. Data was synthetically generated as $\mathbf{y} = \mathbf{X}\mathbf{w}$, where $\mathbf{X} \in \mathbb{R}^{n \times d}$ with i.i.d $\mathcal{N}(0, 1)$ entries and $d = 1000, n = 200$. A held out test dataset of 100 samples was also generated. Other interpolating solutions were computed as $\hat{\mathbf{w}} = \mathbf{X}^\dagger \mathbf{y} + (\mathbf{I} - \mathbf{X}^\dagger \mathbf{X})\mathbf{v}$ and the norm of $\mathbf{v}$ was varied to obtain the plot. Train MSE is 0 for all interpolants, but test MSE increases as $\|\mathbf{v}\|$ increases, with $\mathbf{w}^\dagger$ having the best performance. CV_{100} stability also increases as $\|\mathbf{v}\|$ increases, with the minimum norm interpolant having the best stability. These plots represent results averaged over 100 trials.

obtain RBF and polynomial kernel matrices with different singular values, we varied the size of the training dataset from 10 to 200 (Figs. 2(a) and 2(b)). In both cases, we observe that the log test MSE of the minimum norm interpolating solution is correlated with the log of the norm of the pseudoinverse of the empirical kernel matrix. This confirms our observation that the numerical stability and statistical stability of a kernel least squares problem are related through the smallest singular value of the kernel matrix.



(a) RBF Kernel



(b) Polynomial Kernel

Fig. 2. Plot of log test mean squared error versus $\log \frac{1}{\sigma_{\min}(K)}$ for a kernel least squares problem using (a) An RBF kernel with $\sigma = 5$ and (b) A polynomial kernel of degree 2. We synthetically generated a training dataset $\mathbf{X} \in \mathbb{R}^{n \times d}$ ($d = 20, n = 200$) from i.i.d $\mathcal{N}(0, 1)$ entries and $y_i = \exp(-2\|\mathbf{x}_i\|)$, as well as a held out test dataset of 100 samples. In both (a) and (b), we vary the size of the training dataset from 10 to 200 to obtain kernel matrices with different singular values. These plots represent results averaged over 100 trials.

We note that our results *do not* predict a double descent curve in the smallest singular value for kernels that are not linear dot product kernels [22]. In the case of linear dot product kernels, since the spectra of $\mathbf{X}^\top \mathbf{X}$ and $\mathbf{X} \mathbf{X}^\top$ are the same, one can expect a double descent curve for the smallest singular value of the kernel matrix. This is not true for more general kernels. The expected loss should therefore also not show a double descent, except in the case of the linear kernel. This is also what we find empirically (see Figs. 1(a) and 2(a)).

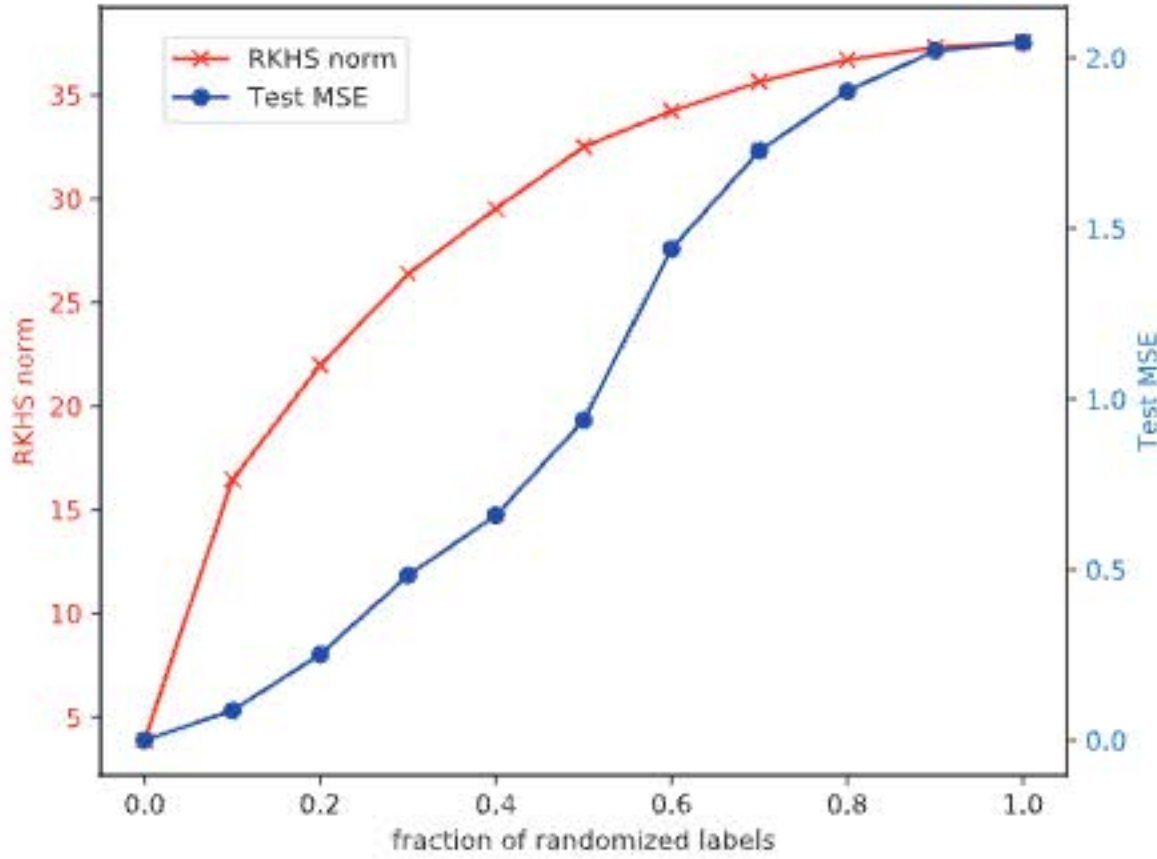


Fig. 3. Effect of randomizing labels on the solution norm and test mean squared error. This plot was generated from a binary classification experiment with an RBF kernel, where the data was drawn from $\mathcal{N}(\pm 3 \times \mathbf{1}_{20}, \mathbf{I}_{20})$ with 200 training samples and 100 test samples.

Test Error and Norm of the Solution: In order to show that the norm of the minimum norm solution to the kernel least squares problem also governs the stability and hence test error (Remark 6.2), we performed an experiment on binary classification in which the fraction of random labels assigned was varied in order to increase the noise level of the problem. We generated data from two gaussian distributions with different means, i.e. from $\mathcal{N}(\pm 3 \times \mathbf{1}_{20}, \mathbf{I}_{20})$, and trained an RBF kernel on 200 training samples and observed the test error on a held out set of 100 samples. As we can see in Fig. 3 the norm of the interpolating solution (red) and the test mean squared error (blue) both increase as the label noise increases. An intuitive explanation of the reason the norm grows is that the pseudoinverse of the data operator $\mathbf{X}^\dagger$ is effectively a high-pass filter that amplifies high-frequencies (more noise) in $\mathbf{y}$, and increases the norm of the minimum norm solution. We also expect the test error to grow as the label noise in the problem increases, since the minimum achievable error is at least the label noise. Our results on the stability and hence test error of the minimum norm solution to the kernel least squares problem also capture this phenomenon.

8. Conclusions

In summary, minimizing a bound on cross validation stability minimizes the expected error in both the classical and the modern regime of ERM. In the classical regime ($d < n$, d large but fixed), CV_{100} stability implies generalization and consistency (for $n \rightarrow \infty$). In the modern regime ($d > n$), as described in this paper,

optimizing CV_{loo} stability selects the minimum norm interpolating solution to the kernel least squares problem which has the best generalization performance.

The main contribution of this paper is in characterizing the stability of (possibly interpolating) solutions to the kernel least squares problem, in particular deriving excess risk bounds via a stability argument. In the process, we show that among all the interpolating solutions, the one with minimum norm also minimizes a bound on stability. Since the excess risk bounds of the minimum norm interpolating solution depend on the minimum singular value of the kernel matrix which is closely related to the condition number, we establish here a link between *numerical and statistical* stability. This also holds for solutions computed by gradient descent, since gradient descent converges to minimum norm solutions in the case of “linear” kernel methods. Our approach is simple and combines basic stability results with matrix inequalities. It is our hope that similar results may be established for deep networks, in particular with respect to minimum norm solutions being the most stable.

Appendix A. Excess Risk, Generalization and Stability

We use the same notation as introduced in Sec. 2 for the various quantities considered in this section. That is in the supervised learning setup $V(f, z)$ is the loss incurred by hypothesis f on the sample z and $I[f] = \mathbb{E}_z[V(f, z)]$ is the expected error of hypothesis f . Since we are interested in different forms of stability, we will consider learning problems over the original training set $S = \{z_1, z_2, \dots, z_n\}$, the leave one out training set $S_{-i} = \{z_1, \dots, z_{i-1}, z_{i+1}, \dots, z_n\}$, and the replace one training set $(S_{-i}, z) = \{z_1, \dots, z_{i-1}, z_{i+1}, \dots, z_n, z\}$

A.1. Replace one and leave one out algorithmic stability

Similar to the definition of expected CV_{loo} stability in Eq. (3.1) of the main paper, we say an algorithm is cross validation *replace one* stable (in expectation), denoted as CV_{ro} , if there exists $\beta_{\text{ro}} > 0$ such that

$$\mathbb{E}_{S,z}[V(f_S, z) - V(f_{(S_{-i}, z)}, z)] \leq \beta_{\text{ro}}.$$

We can strengthen the above stability definition by introducing the notion of replace one algorithmic stability (in expectation) [7]. There exists $\alpha_{\text{ro}} > 0$ such that for all $i = 1, \dots, n$,

$$\mathbb{E}_{S,z}[\|f_S - f_{(S_{-i}, z)}\|_\infty] \leq \alpha_{\text{ro}}.$$

We make two observations.

First, if the loss is Lipschitz, that is if there exists $C_V > 0$ such that for all $f, f' \in \mathcal{H}$,

$$\|V(f, z) - V(f', z)\| \leq C_V \|f - f'\|,$$

then replace one algorithmic stability implies CV_{ro} stability with $\beta_{\text{ro}} = C_V \alpha_{\text{ro}}$. Moreover, the same result holds if the loss is locally Lipschitz and there exists

$R > 0$, such that $\|f_S\|_\infty \leq R$ almost surely. In this latter case, the Lipschitz constant will depend on R . Later, we illustrate this situation for the square loss.

Second, we have for all $i = 1, \dots, n$, S and z ,

$$\mathbb{E}_{S,z}[\|f_S - f_{(S_{-i}, z)}\|_\infty] \leq \mathbb{E}_{S,z}[\|f_S - f_{S_{-i}}\|_\infty] + \mathbb{E}_{S,z}[\|f_{(S_{-i}, z)} - f_{S_{-i}}\|_\infty].$$

This observation motivates the notion of leave one out algorithmic stability (in expectation) [7]:

$$\mathbb{E}_{S,z}[\|f_S - f_{S_{-i}}\|_\infty] \leq \alpha_{\text{loo}}.$$

Clearly, leave one out algorithmic stability implies replace one algorithmic stability with $\alpha_{\text{ro}} = 2\alpha_{\text{loo}}$ and it implies also CV_{ro} stability with $\beta_{\text{ro}} = 2C_V\alpha_{\text{loo}}$.

A.2. Excess risk and CV_{loo} , CV_{ro} stability

We recall the statement of Lemma 3.1 in Sec. 3 that bounds the excess risk using the CV_{loo} stability of a solution.

Lemma A.1 (Excess Risk and CV_{loo} Stability). *For all $i = 1, \dots, n$,*

$$\mathbb{E}_S \left[I[f_{S_{-i}}] - \inf_{f \in \mathcal{H}} I[f] \right] \leq \mathbb{E}_S [V(f_{S_{-i}}, z_i) - V(f_S, z_i)]. \quad (\text{A.1})$$

In this section, two properties of ERM are useful, namely symmetry, and a form of unbiasedness.

Symmetry. A key property of ERM is that it is *symmetric* with respect to the data set S , meaning that it does not depend on the order of the data in S .

A second property relates the expected ERM with the minimum of expected risk.

ERM Bias. The following inequality holds:

$$\mathbb{E}[[I_S[f_S]] - \min_{f \in \mathcal{H}} I[f]] \leq 0. \quad (\text{A.2})$$

To see this, note that

$$I_S[f_S] \leq I_S[f]$$

for all $f \in \mathcal{H}$ by definition of ERM, so that taking the expectation of both sides

$$\mathbb{E}_S[I_S[f_S]] \leq \mathbb{E}_S[I_S[f]] = I[f]$$

for all $f \in \mathcal{H}$. This implies

$$\mathbb{E}_S[I_S[f_S]] \leq \min_{f \in \mathcal{H}} I[f]$$

and hence (A.2) holds.

Remark A.1. Note that the same argument gives more generally that

$$\mathbb{E} \left[\inf_{f \in \mathcal{H}} [I_S[f]] \right] - \inf_{f \in \mathcal{H}} I[f] \leq 0. \quad (\text{A.3})$$

Given the above premise, the proof of Lemma 3.1 is simple.

Proof of Lemma 3.1. Adding and subtracting $\mathbb{E}_S[I_S[f_S]]$ from the expected excess risk we have that

$$\mathbb{E}_S \left[I[f_{S_{-i}}] - \min_{f \in \mathcal{H}} I[f] \right] = \mathbb{E}_S \left[I[f_{S_{-i}}] - I_S[f_S] + I_S[f_S] - \min_{f \in \mathcal{H}} I[f] \right] \quad (\text{A.4})$$

and since $\mathbb{E}_S[I_S[f_S]] - \min_{f \in \mathcal{H}} I[f]$ is less or equal than zero, see (A.3), then

$$\mathbb{E}_S \left[I[f_{S_{-i}}] - \min_{f \in \mathcal{H}} I[f] \right] \leq \mathbb{E}_S[I[f_{S_{-i}}] - I_S[f_S]]. \quad (\text{A.5})$$

Moreover, for all $i = 1, \dots, n$,

$$\mathbb{E}_S[I[f_{S_{-i}}]] = \mathbb{E}_S[\mathbb{E}_{z_i} V(f_{S_{-i}}, z_i)] = \mathbb{E}_S[V(f_{S_{-i}}, z_i)]$$

and

$$\mathbb{E}_S[I_S[f_S]] = \frac{1}{n} \sum_{i=1}^n \mathbb{E}_S[V(f_S, z_i)] = \mathbb{E}_S[V(f_S, z_i)].$$

Plugging these last two expressions in (A.5) and in (A.4) leads to (3.2). $\square$

We can prove a similar result relating excess risk with CV_{ro} stability.

Lemma A.2 (Excess Risk and CV_{ro} Stability). *Given the above definitions, the following inequality holds for all $i = 1, \dots, n$:*

$$\begin{aligned} \mathbb{E}_S \left[I[f_S] - \inf_{f \in \mathcal{H}} I[f] \right] &\leq \mathbb{E}_S[I[f_S] - I_S[f_S]] \\ &= \mathbb{E}_{S,z}[V(f_S, z) - V(f_{(S_{-i}, z)}, z)]. \end{aligned} \quad (\text{A.6})$$

Proof. The first inequality is clear from adding and subtracting $I_S[f_S]$ from the expected risk $I[f_S]$ we have that

$$\mathbb{E}_S \left[I[f_S] - \min_{f \in \mathcal{H}} I[f] \right] = \mathbb{E}_S \left[I[f_S] - I_S[f_S] + I_S[f_S] - \min_{f \in \mathcal{H}} I[f] \right]$$

and recalling (A.3). The main step in the proof is showing that for all $i = 1, \dots, n$,

$$\mathbb{E}[I_S[f_S]] = \mathbb{E}[V(f_{(S_{-i}, z)}, z)] \quad (\text{A.7})$$

to be compared with the trivial equality, $\mathbb{E}[I_S[f_S]] = \mathbb{E}[V(f_S, z_i)]$. To prove Eq. (A.7), we have for all $i = 1, \dots, n$,

$$\begin{aligned} \mathbb{E}_S[I_S[f_S]] &= \mathbb{E}_{S,z} \left[\frac{1}{n} \sum_{i=1}^n V(f_S, z_i) \right] = \frac{1}{n} \sum_{i=1}^n \mathbb{E}_{S,z}[V(f_{(S_{-i}, z)}, z)] \\ &= \mathbb{E}_{S,z}[V(f_{(S_{-i}, z)}, z)], \end{aligned}$$

where we used the fact that by the symmetry of the algorithm $\mathbb{E}_{S,z}[V(f_{(S_{-i}, z)}, z)]$ is the same for all $i = 1, \dots, n$. The proof is concluded noting that $\mathbb{E}_S[I[f_S]] = \mathbb{E}_{S,z}[V(f_S, z)]$. $\square$

A.3. Discussion on stability and generalization

Below we discuss some more aspects of stability and its connection to other quantities in statistical learning theory.

Remark A.2 (CV_{loo} Stability in Expectation and in Probability). In [19], CV_{loo} stability is defined in probability, that is there exists $\beta_{CV}^P > 0$, $0 < \delta_{CV}^P \leq 1$ such that

$$\mathbb{P}_S\{|V(f_{S_{-i}}, z_i) - V(f_S, z_i)| \geq \beta_{CV}^P\} \leq \delta_{CV}^P.$$

Note that the absolute value is not needed for ERM since almost positivity holds [19], that is $V(f_{S_{-i}}, z_i) - V(f_S, z_i) > 0$. Then CV_{loo} stability in probability and in expectation are clearly related and indeed equivalent for bounded loss functions. CV_{loo} stability in expectation (3.1) is what we study in the following sections.

Remark A.3 (Connection to Uniform Stability and Other Notions of Stability). Uniform stability, introduced by [7], corresponds in our notation to the assumption that there exists $\beta_u > 0$ such that for all $i = 1, \dots, n$, $\sup_{z \in Z} |V(f_{S_{-i}}, z) - V(f_S, z)| \leq \beta_u$. Clearly, this is a strong notion implying most other definitions of stability. We note that there are a number of different notions of stability. We refer the interested reader to [11, 19].

Remark A.4 (CV_{loo} Stability and Learnability). A natural question is to which extent suitable notions of stability are not only sufficient but also necessary for controlling the excess risk of ERM. Classically, the latter is characterized in terms of a uniform version of the law of large numbers, which itself can be characterized in terms of suitable complexity measures of the hypothesis class. Uniform stability is too strong to characterize consistency while CV_{loo} stability turns out to provide a suitably weak definition as shown in [19], see also [11, 19]. Indeed, a main result in [19] shows that CV_{loo} stability is *equivalent to consistency of ERM*.

Theorem A.1 ([19]). *For ERM and bounded loss functions, CV_{loo} stability in probability with β_{CV}^P converging to zero for $n \rightarrow \infty$ is equivalent to consistency and generalization of ERM.*

Remark A.5 (CV_{loo} Stability and In-Sample/Out-of-Sample Error). Let $(S, z) = \{z_1, \dots, z_n, z\}$ (z is a data point drawn according to the same distribution) and the corresponding ERM solution $f_{(S, z)}$, then (3.2) can be equivalently written as

$$\mathbb{E}_S \left[I[f_S] - \inf_{f \in \mathcal{F}} I[f] \right] \leq \mathbb{E}_{S, z} [V(f_S, z) - V(f_{(S, z)}, z)].$$

Thus, CV_{loo} stability measures how much the loss changes when we test on a point that is present in the training set and absent from it. In this view, it can be seen as an average measure of the difference between in-sample and out-of-sample error.

Remark A.6 (CV_{loo} Stability and Generalization). A common error measure is the (expected) generalization gap $\mathbb{E}_S[I[f_S] - I_S[f_S]]$. For non-ERM algorithms, CV_{loo} stability by itself is not sufficient to control this term, and further conditions are needed [19], since

$$\mathbb{E}_S[I[f_S] - I_S[f_S]] = \mathbb{E}_S[I[f_S] - I_S[f_{S_{-i}}]] + \mathbb{E}_S[I_S[f_{S_{-i}}] - I_S[f_S]].$$

The second term becomes for all $i = 1, \dots, n$,

$$\begin{aligned} \mathbb{E}_S[I_S[f_{S_{-i}}] - I_S[f_S]] &= \frac{1}{n} \sum_{i=1}^n \mathbb{E}_S[V(f_{S_{-i}}, z_i) - V(f_S, z_i)] \\ &= \mathbb{E}_S[V(f_{S_{-i}}, z_i) - V(f_S, z_i)] \end{aligned}$$

and hence is controlled by CV stability. The first term is called expected leave one out error in [19] and is controlled in ERM as $n \rightarrow \infty$, see Theorem A.1 above.

Appendix B. Generalized Inverse of a Perturbed Operator

In this section, we consider a linear operator perturbed by a rank-one operator, i.e. $\mathbf{M} = \mathbf{A} + \mathbf{c}\mathbf{d}^\top$, where $\mathbf{M}, \mathbf{A}: U \rightarrow V$, $\mathbf{d} \in U$, $\mathbf{c} \in V$. Here U, V are inner product vector spaces over the field $\mathbb{R}$.

Theorem B.1 ([17, Theorem 6]). Let $\mathbf{A}^\dagger$ be the pseudoinverse of $\mathbf{A}$, and define $\mathbf{k} = \mathbf{A}^\dagger \mathbf{c}$, $\mathbf{h} = \mathbf{d}^\top \mathbf{A}^\dagger$. Let us consider the case where $\mathbf{c} \in \text{range}(\mathbf{A})$, $\mathbf{d} \in \text{range}(\mathbf{A}^\top)$ and $\beta = 1 + \mathbf{d}^\top \mathbf{A}^\dagger \mathbf{c} = 0$. Then the pseudoinverse $\mathbf{M}^\dagger$ of the perturbed operator is given by

$$\mathbf{M}^\dagger = \mathbf{A}^\dagger - \mathbf{k}\mathbf{k}^\dagger \mathbf{A}^\dagger - \mathbf{A}^\dagger \mathbf{h}^\dagger \mathbf{h} + (\mathbf{k}^\dagger \mathbf{A}^\dagger \mathbf{h}^\dagger) \mathbf{k} \mathbf{h}. \quad (\text{B.1})$$

Proof. We will reproduce the proof of [17, Theorem 6] here, with the only difference being that instead of matrices, we have linear operators.

Consider the operators $\mathbf{A}\mathbf{A}^\dagger - \mathbf{h}^\dagger \mathbf{h}$, $\mathbf{A}^\dagger \mathbf{A} - \mathbf{k}\mathbf{k}^\dagger$. These are both orthogonal projectors, since they are symmetric and idempotent. This can be checked quite easily, using the facts that $\mathbf{A}\mathbf{A}^\dagger \mathbf{h}^\dagger = \mathbf{h}^\dagger$, $\mathbf{h}\mathbf{A}\mathbf{A}^\dagger = \mathbf{h}$, $\mathbf{A}^\dagger \mathbf{A} \mathbf{k} = \mathbf{k}$ and $\mathbf{k}^\dagger \mathbf{A}^\dagger \mathbf{A} = \mathbf{k}^\dagger$. Both of these operators have their rank equal to $\text{rank}(\mathbf{A}) - 1$. This is also the case for the operator $\mathbf{M}$, from [17, Lemma 1].

Hence, we have $\text{rank}(\mathbf{M}) = \text{rank}(\mathbf{A}\mathbf{A}^\dagger - \mathbf{h}^\dagger \mathbf{h}) = \text{rank}(\mathbf{A}^\dagger \mathbf{A} - \mathbf{k}\mathbf{k}^\dagger)$.

With the facts that $\mathbf{A}\mathbf{A}^\dagger \mathbf{c} = \mathbf{c}$, $\mathbf{h}\mathbf{c} = -1$ and $\mathbf{h}\mathbf{A} = \mathbf{d}^\top$, we have that

$$(\mathbf{A}\mathbf{A}^\dagger - \mathbf{h}^\dagger \mathbf{h})\mathbf{M} = \mathbf{M}.$$

This means that $\text{range}(\mathbf{M}) \subset \text{range}(\mathbf{A}\mathbf{A}^\dagger - \mathbf{h}^\dagger \mathbf{h})$.

Likewise, with the facts that $\mathbf{d}^\top \mathbf{A}^\dagger \mathbf{A} = \mathbf{d}^\top$, $\mathbf{d}^\top \mathbf{k} = -1$ and $\mathbf{A}\mathbf{k} = \mathbf{c}$, we have that

$$\mathbf{M}(\mathbf{A}^\dagger \mathbf{A} - \mathbf{k}\mathbf{k}^\dagger) = \mathbf{M}$$

and hence $\text{range}(\mathbf{M}^\top) \subset \text{range}(\mathbf{A}^\dagger \mathbf{A} - \mathbf{k}\mathbf{k}^\dagger)$. putting these together, we have that

$$\begin{aligned} \mathbf{M}\mathbf{M}^\dagger &= \mathbf{A}\mathbf{A}^\dagger - \mathbf{h}^\dagger \mathbf{h}, \\ \mathbf{M}^\dagger \mathbf{M} &= \mathbf{A}^\dagger \mathbf{A} - \mathbf{k}\mathbf{k}^\dagger. \end{aligned} \tag{B.2}$$

If $\mathbf{X}$ is the right-hand side of [B.1] we can use $\mathbf{h}\mathbf{A}\mathbf{A}^\dagger = \mathbf{h}$ and the above equation to obtain $\mathbf{X}\mathbf{M}\mathbf{M}^\dagger = \mathbf{X}$. We can also use the above equation, $\mathbf{k}^\dagger \mathbf{A}^\dagger \mathbf{A} = \mathbf{k}^\dagger$, $\mathbf{h}\mathbf{A} = \mathbf{d}^\top$ and $\mathbf{h}\mathbf{c} = -1$ to obtain $\mathbf{X}\mathbf{M} = \mathbf{M}^\dagger \mathbf{M}$. Since this means that $\mathbf{X}$ satisfies the two conditions of [17, Lemma 2], we have shown that $\mathbf{M}^\dagger = \mathbf{X}$. $\square$

Lemma B.1 ([17, Lemma 1]). *Let $\mathbf{A}^\dagger$ be the pseudoinverse of $\mathbf{A}$, and define $\mathbf{u} = (\mathbf{I} - \mathbf{A}\mathbf{A}^\dagger)\mathbf{c}$, $\mathbf{v} = \mathbf{d}^\top (\mathbf{I} - \mathbf{A}^\dagger \mathbf{A})$ and $\beta = 1 + \mathbf{d}^\top \mathbf{A}^\dagger \mathbf{c} = 0$. Then the rank of the perturbed operator $\mathbf{M} = \mathbf{A} + \mathbf{c}\mathbf{d}^\top$ is given by*

$$\text{rank}(\mathbf{A} + \mathbf{c}\mathbf{d}^\top) = \text{rank} \begin{bmatrix} \mathbf{A} & \mathbf{u} \\ \mathbf{v} & -\beta \end{bmatrix} - 1.$$

Lemma B.2 ([17, Lemma 2]). *If $\mathbf{X}$ and $\mathbf{M}$ are operators such that $\mathbf{X}\mathbf{M}\mathbf{M}^\dagger = \mathbf{X}$ and $\mathbf{M}^\dagger \mathbf{M} = \mathbf{X}\mathbf{M}$, then $\mathbf{X} = \mathbf{M}^\dagger$.*

References

- [1] S. Arora, S. S. Du, W. Hu, Z. Y. Li and R. Wang, Fine-grained analysis of optimization and generalization for overparameterized two-layer neural networks, preprint (2019), arXiv:1901.08584.
- [2] J. K. Baksalary, O. M. Baksalary and G. Trenkler, A revisitation of formulae for the Moore–Penrose inverse of modified matrices, *Linear Algebra Appl.* **372** (2003) 207–224.
- [3] P. L. Bartlett, P. M. Long, G. Lugosi and A. Tsigler, Benign overfitting in linear regression, preprint (2019), arXiv:1906.11300.
- [4] M. Belkin, D. Hsu, S. Ma and S. Mandal, Reconciling modern machine-learning practice and the classical bias–variance trade-off, *Proc. Natl. Acad. Sci. USA* **116**(32) (2019) 15849–15854, doi:10.1073/pnas.1903070116.
- [5] M. Belkin, S. Ma and S. Mandal, To understand deep learning we need to understand kernel learning, in *Proc. 35th Int. Conf. Machine Learning*, Vol. 80, eds. J. Dy and A. Krause (PMLR, 2018), pp. 541–549, <https://proceedings.mlr.press/v80/belkin18a.html>.
- [6] S. Boucheron, O. Bousquet and G. Lugosi, Theory of classification: A survey of some recent advances, *ESAIM Probab. Stat.* **9** (2005) 323–375.
- [7] O. Bousquet and A. Elisseeff, Stability and generalization, *J. Mach. Learn. Res.* **2** (2002) 499–526.
- [8] P. Bühlmann and S. Van De Geer, *Statistics for High-Dimensional Data: Methods, Theory and Applications* (Springer Science & Business Media, 2011).

- [9] N. El Karoui, The spectrum of kernel random matrices, preprint (2010), arXiv:1001.0492.
- [10] T. Hastie, A. Montanari, S. Rosset and R. J. Tibshirani, Surprises in high-dimensional ridgeless least squares interpolation, preprint (2019), arXiv:1903.08560.
- [11] S. Kutin and P. Niyogi, Almost-everywhere algorithmic stability and generalization error, Technical Report TR-2002-03, University of Chicago, 2002.
- [12] T. Liang, A. Rakhlin and X. Zhai, On the risk of minimum-norm interpolants and restricted lower isometry of kernels, preprint (2019), arXiv:1908.10292.
- [13] T. Liang *et al.*, Just interpolate: Kernel “ridgeless” regression can generalize, *Ann. Statist.* **48**(3) (2020) 1329–1347.
- [14] T. Liang and B. Recht, Interpolating classifiers make few mistakes, preprint (2021), arXiv:2101.11815.
- [15] V. A. Marchenko and L. A. Pastur, Distribution of eigenvalues for some sets of random matrices, *Mat. Sb. (N.S.)* **72**(114)(4) (1967) 457–483.
- [16] S. Mei and A. Montanari, The generalization error of random features regression: Precise asymptotics and double descent curve, preprint (2019), arXiv:1908.05355.
- [17] C. Meyer, Generalized inversion of modified matrices, *SIAM J. Appl. Math.* **24** (1973) 315–323.
- [18] C. A. Micchelli, Interpolation of scattered data: Distance matrices and conditionally positive definite functions, *Constr. Approx.* **2** (1986) 11–22.
- [19] S. Mukherjee, P. Niyogi, T. Poggio and R. Rifkin, Learning theory: Stability is sufficient for generalization and necessary and sufficient for consistency of empirical risk minimization, *Adv. Comput. Math.* **25**(1) (2006) 161–193, doi:10.1007/s10444-004-7634-z.
- [20] E. Oravkin and P. Rebeschini, On optimal interpolation in linear regression, in *35th Conf. Neural Information Processing Systems* (NeurIPS 2021), pp. 1–12.
- [21] T. Poggio, *Stable Foundations for Learning* (Center for Brains, Minds and Machines, 2020).
- [22] T. Poggio, G. Kur and A. Banburski, Double descent in the condition number, Technical Report, MIT Center for Brains Minds and Machines, 2019.
- [23] T. Poggio, R. Rifkin, S. Mukherjee and P. Niyogi, General conditions for predictivity in learning theory, *Nature* **428** (2004) 419–422.
- [24] A. Rakhlin and X. Zhai, Consistency of interpolation with Laplace kernels is a high-dimensional phenomenon, preprint (2018), arXiv:1812.11167.
- [25] L. Rosasco and S. Villa, Learning with incremental iterative regularization, In *Advances in Neural Information Processing Systems*, Vol. 28, eds. C. Cortes, N. D. Lawrence, D. D. Lee, M. Sugiyama and R. Garnett (Curran Associates, 2015), pp. 1630–1638, <http://papers.nips.cc/paper/6015-learning-with-incremental-iterative-regularization.pdf>.
- [26] S. Shalev-Shwartz and S. Ben-David, *Understanding Machine Learning: From Theory to Algorithms* (Cambridge University Press, USA, 2014).
- [27] S. Shalev-Shwartz, O. Shamir, N. Srebro and K. Sridharan, Learnability, stability and uniform convergence, *J. Mach. Learn. Res.* **11** (2010) 2635–2670.
- [28] I. Steinwart and A. Christmann, *Support Vector Machines* (Springer Science & Business Media, 2008).
- [29] C. Zhang, S. Bengio, M. Hardt, B. Recht and O. Vinyals, Understanding deep learning requires rethinking generalization, in *5th Int. Conf. Learning Representations, ICLR 2017, Toulon, France, April 24–26, 2017, Conference Track Proc.* (OpenReview.net, 2017), pp. 1–15, <https://openreview.net/forum?id=Sy8gdB9xx>.