

SURPRISES IN HIGH-DIMENSIONAL RIDGELESS LEAST SQUARES INTERPOLATION

BY TREVOR HASTIE^{1,a}, ANDREA MONTANARI^{2,b}, SAHARON ROSSET^{3,c} AND RYAN J. TIBSHIRANI^{4,d}

¹*Department of Statistics and Department of Biomedical Data Science, Stanford University, hastie@stanford.edu*

²*Department of Statistics and Department of Electrical Engineering, Stanford University, montanar@stanford.edu*

³*School of Mathematical Sciences, Tel Aviv University, saharon@post.tau.ac.il*

⁴*Department of Statistics and Department of Machine Learning, Carnegie Mellon University, ryantibs@cmu.edu*

Interpolators—estimators that achieve zero training error—have attracted growing attention in machine learning, mainly because state-of-the-art neural networks appear to be models of this type. In this paper, we study minimum ℓ_2 norm (“ridgeless”) interpolation least squares regression, focusing on the high-dimensional regime in which the number of unknown parameters p is of the same order as the number of samples n . We consider two different models for the feature distribution: a linear model, where the feature vectors $x_i \in \mathbb{R}^p$ are obtained by applying a linear transform to a vector of i.i.d. entries, $x_i = \Sigma^{1/2} z_i$ (with $z_i \in \mathbb{R}^p$); and a nonlinear model, where the feature vectors are obtained by passing the input through a random one-layer neural network, $x_i = \varphi(Wz_i)$ (with $z_i \in \mathbb{R}^d$, $W \in \mathbb{R}^{p \times d}$ a matrix of i.i.d. entries, and φ an activation function acting componentwise on Wz_i). We recover—in a precise quantitative way—several phenomena that have been observed in large-scale neural networks and kernel machines, including the “double descent” behavior of the prediction risk, and the potential benefits of overparametrization.

1. Introduction. Modern deep learning models involve a huge number of parameters. In many applications, current practice suggests that we should design the network to be sufficiently complex so that the model (as trained, typically, by gradient descent) interpolates the data, that is, achieves zero training error. Indeed, in a thought-provoking experiment, Zhang et al. [71] showed that state-of-the-art deep neural network architectures are complex enough that they can be trained to interpolate the data even when the actual labels are replaced by entirely random ones.

Despite their enormous complexity, deep neural networks are frequently observed to generalize well in practice. At first sight, this seems to defy conventional statistical wisdom: interpolation (vanishing training error) is commonly taken to be a proxy for poor generalization (large gap between training and test error), and hence large test error. In an insightful series of papers, Belkin et al. [8, 10, 11] pointed out that these concepts are in general distinct, and interpolation does not contradict generalization. For example, recent work has investigated interpolation—via kernel ridge regression—in reproducing kernel Hilbert spaces [30, 47]. While in low dimension a positive regularization is needed to achieve good interpolation, in certain high-dimensional settings interpolation can be nearly optimal.

In this paper, we investigate these phenomena in the context of simple linear models. We assume to be given i.i.d. data (y_i, x_i) , $i \leq n$, with $x_i \in \mathbb{R}^p$ a feature vector and $y_i \in \mathbb{R}$ a

Received June 2019; revised August 2021.

MSC2020 subject classifications. Primary 62J05, 62J07; secondary 62J02, 62F12.

Key words and phrases. Regression, interpolation, overparametrization, ridge regression, random matrix theory.

response variable. These are distributed according to the model (see Section 2 for further definitions)

$$(1) \quad (x_i, \epsilon_i) \sim P_x \times P_\epsilon, \quad i = 1, \dots, n,$$

$$(2) \quad y_i = x_i^T \beta + \epsilon_i, \quad i = 1, \dots, n,$$

where P_x is a distribution on $\mathbb{R}^p$ such that $\mathbb{E}(x_i) = 0$, $\text{Cov}(x_i) = \Sigma$, and P_ϵ is a distribution on $\mathbb{R}$ such that $\mathbb{E}(\epsilon_i) = 0$, $\text{Var}(\epsilon_i) = \sigma^2$.

We estimate β by linear regression. Since our focus is on the overparametrized regime $p > n$, the usual least square objective does not have a unique minimizer, and needs to be regularized. We consider two approaches: min-norm regression, which estimates β by the least squares solution with minimum ℓ_2 norm; and ridge regression, which penalizes a coefficients vector β by its ℓ_2 norm square $\|\beta\|_2^2$. We denote these estimates by $\hat{\beta}$ and $\hat{\beta}_\lambda$ (λ being the regularization parameter), and note that $\lim_{\lambda \rightarrow 0} \hat{\beta}_\lambda = \hat{\beta}$. If the design matrix has full row rank, which is generically the case for $p > n$, the min-norm estimator is an interpolator, namely $x_i^T \hat{\beta} = y_i$ for all $i \leq n$. In order to evaluate these methods, we will study the prediction risk at a new (unseen) test point (y_0, x_0) .

We study the model (2) in the proportional regime $p \asymp n$, with a special focus on the overparametrized case $p > n$. Our main contribution is to show that, by considering different choices of the features distribution P_x , we can reproduce a number of statistically interesting phenomena that have emerged in the context of deep learning.

From a technical perspective, our main results are: Theorems 2 and 5, which assume the linear model $x_i = \Sigma^{1/2} z_i$ with z_i a vector with independent coordinates and Theorem 8, which assumes a nonlinear model $x_i = \varphi(W z_i)$ with $z_i \sim N(0, I_d)$. While the linear model has already attracted significant amount of work (see Section 1.3 for an overview), Theorems 2 and 5 provide a more accurate approximation of the prediction risk in the proportional regime $n \asymp p$, as compared to available results in the literature, and hold in a more general setting.

The prediction risk depends on the geometry of the pair (Σ, β) . We consider a few different choices for this geometry, which are broadly motivated by our objective to understand overparametrized models, and specialize our formulas to these special cases:

1. *Isotropic features.* This is the simplest case, in which $\Sigma = I_p$ and, therefore, as we will see the asymptotic risk depends on β only through its norm $\|\beta\|_2$. This simple model captures some interesting features of overparametrization, but misses others.

We first consider a well-specified case in which $x_i \in \mathbb{R}^p$ and we regress against x_i . We then pass to a misspecified case, in which the model (2) holds for covariates $x_i \in \mathbb{R}^{p+q}$, but we regress only against the first p covariates.

2. *Latent space features.* In the overparametrized regime, it is natural to assume that both the covariates x_i , and the coefficients vector β lie close to a low-dimensional subspace. In order to model this property, we assume $\Sigma = W W^T + I$, with $W \in \mathbb{R}^{p \times d}$, $d \ll p$ and β lies in the span of the columns of W . Interestingly, this model reproduces many phenomena observed in more complex nonlinear models, and has a more direct connection to neural networks.

3. *Nonlinear model.* In all of the previous cases, the distribution of x_i is of the form $x_i = \Sigma^{1/2} z_i$ where z_i is a vector with independent coordinates. In order to test the generality of our results, we consider a model in which x_i is obtained by passing $z_i \sim N(0, I_d)$ through a one-layer neural net with random first layer weights, namely $x_i = \varphi(W z_i)$, for $W \in \mathbb{R}^{p \times d}$.

We will summarize our results for these four examples in the next subsection.

A skeptical reader might ask what linear models have to do with neural networks. We emphasize that linear models provide more than a simple analogy, and a recent line of work

outlines a concrete connection between the two settings [4, 19, 24, 25, 39]. We will discuss this connection in Section 1.2.

1.1. Summary of results. As mentioned above, we analyze the out-of-sample prediction risk of the minimum ℓ_2 norm (or min-norm, for short) least squares estimator, and of ridge-regularized least squares.

We denote by $\gamma_n := p/n \in (0, \infty)$ the overparametrization ratio. When $\gamma < 1$, we call the problem *underparametrized*, and when $\gamma > 1$, we call it *overparametrized*. Our most general results for the linear model (Theorem 2 and 5) apply to a nonasymptotic setting in which n, p are finite, and provide a deterministic approximation of the risk with error bounds that are uniform in the distribution of the data. We use these general results to derive asymptotic formulas in the limit in which both p and n diverge with $\gamma_n = p/n \rightarrow \gamma$. (We will drop the subscript from γ_n whenever this is not cause for confusion.)

We assume the model (2) and denote by $\text{SNR} = \|\beta\|_2^2/\sigma^2$ the signal-to-noise ratio. We refer to Figure 1 for supporting plots of the asymptotic risk curves for different cases of interest.

Our main results are twofold: (i) We show that by suitable choices of β, Σ , we can easily construct scenarios in which the minimum of the risk is achieved in the overparametrized regime $p > n$; (ii) We show that these findings are robust to the details of the distribution of (y_i, x_i) .

As a preliminary remark, note that in the underparametrized regime ($\gamma < 1$), the min-norm estimator coincides with the standard least squares estimator. Its risk is purely variance (there

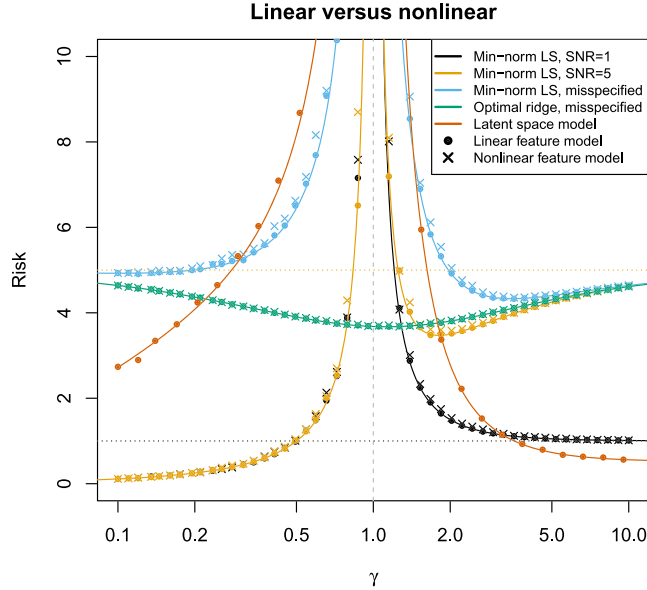


FIG. 1. Asymptotic risk curves for the linear feature model, as a function of the limiting aspect ratio γ . Black and yellow: risks for min-norm least squares in the isotropic well-specified model, for $\text{SNR} = 1$ and $\text{SNR} = 5$, respectively. These two match for $\gamma < 1$ but differ for $\gamma > 1$. The null risks for $\text{SNR} = 1$ and $\text{SNR} = 5$ are marked by the dotted black and yellow lines, respectively. Light blue: risk for a misspecified model with significant approximation bias ($a = 1.5$ in (27)), when $\text{SNR} = 5$. Green: optimally-tuned (equivalently, CV-tuned) ridge regression, in the same misspecified setup as for the light blue. Red: latent space model of Section 5.2, with $r = 7$, $\sigma = 0$. The points denote finite-sample risks, with $n = 200$, $p = \lfloor \gamma n \rfloor$, across various values of γ . Meanwhile, the “x” points mark finite-sample risks for a nonlinear feature model, with $n = 200$, $p = \lfloor \gamma n \rfloor$, $d = 100$ and $X = \varphi(ZW^T)$, where Z has i.i.d. $N(0, 1)$ entries, W has i.i.d. $N(0, 1/d)$ entries and $\varphi(t) = a(|t| - b)$ is a “purely nonlinear” activation function, for constants a, b . Theorem 8 predicts that this nonlinear risk should converge to the linear risk with p features (regardless of d).

is no bias), and does not depend on β , Σ (see Proposition 2). Interestingly, the asymptotic risk diverges as we approach the interpolation boundary (as $\gamma \rightarrow 1$).

In contrast, in the overparametrized regime ($\gamma > 1$), the risk is composed of both bias and variance,¹ and generally depends on β , Σ (see Theorem 2).

We next highlight some concrete results for the four models discussed in the previous section (unless explicitly said, we refer to the min-norm estimator).

Isotropic features. The asymptotic risk depends on the coefficients vector only through its norm $\|\beta\|_2^2$ or, up to a scaling, on $\text{SNR} = \|\beta\|_2^2/\sigma^2$.

1. If the model is *well specified*, we observe two different behaviors. For $\text{SNR} \leq 1$, the risk is decreasing for $\gamma \in (1, \infty)$. For $\text{SNR} > 1$, the risk has a *local* minimum on $\gamma \in (1, \infty)$.

In either case, the risk approaches the null risk as $\gamma \rightarrow \infty$, and achieves its global minimum in the underparametrized regime (see Section 3.2).

2. If the model is *misspecified*, when $\text{SNR} > 1$, the risk can attain its *global* minimum in the overparametrized regime $\gamma \in (1, \infty)$ (when there is strong enough approximation bias, see Section 5.1.3). However, the risk is again increasing for γ large enough.

3. Optimally-tuned ridge regression uses a nonvanishing regularization $\lambda > 0$, and dominates the min-norm least squares estimator in risk, across all values of γ and SNR , both the well-specified and misspecified settings. For a misspecified model, optimally-tuned ridge regression attains its global minimum around $\gamma = 1$ (see Section 6).

4. Optimal tuning of the ridge penalty can be achieved by leave-one-out cross-validation (see Theorem 7).

Anisotropic features. In this case, $\Sigma \neq I$ and the risk depends on the geometry of (Σ, β) , and in particular on how β aligns with the eigenvectors of Σ .

1. If the coefficients vector is equidistributed along the eigenvectors of Σ , the behavior is qualitatively similar to the isotropic case. This situation arises, for instance, if β is itself random with a spherical prior.

2. If β is aligned with the top eigenvectors of Σ , the situation is qualitatively different. As an example we obtain an explicit formula for the asymptotic risk in the latent space model discussed above; see the red line of Figure 1 (and Figure 5) for an illustration. We find that, for natural choices of the model parameters, the risk is monotone decreasing in the overparametrized regime, and reaches its global minimum as $\gamma \rightarrow \infty$. This qualitative behavior matches the one observed for neural networks (see Section 5.2).

3. For the latent space model, we observe that, at large overparametrization, the minimum error is achieved as $\lambda \rightarrow 0$, that is, by min-norm interpolators (see Section 6.2, and Section 1.3 for related work).

Nonlinear model. Finally, we consider a nonlinear model in which $x_i = \varphi(Wz_i)$ where φ is a nonlinear activation function applied componentwise, $W \in \mathbb{R}^{p \times d}$ and $z_i \sim N(0, I_d)$.

1. We first consider the case of a purely nonlinear activation. We compute the limiting risk of min-norm regression, and show that this matches the one for Gaussian $x_i \sim N(0, I_d)$ (see Theorem 8). This is illustrated by the “x” symbols in Figure 1.

2. We then compute the limit of the variance component of the risk for more general activations φ . We show that this depends on the activation function only through the size

¹Note that in the overparametrized regime the bias is nonvanishing even in the interpolation limit $\lambda \rightarrow 0$. The reason is that the set of interpolators is an affine space of dimension $p - n$, and the min-norm criterion selects one specific interpolator, whose mean has—in general—norm smaller than β .

of its linear component. Further, the resulting variance turns out to coincide asymptotically with the variance in the linear model $x_i = \Sigma^{1/2} \tilde{z}_i$, if we take $\Sigma = (1 - c_1)I_p + c_1 WW^T$ for a certain constant c_1 . (see Theorem 9).

These results confirm that the results established for the case $x_i = \Sigma^{1/2} z_i$ with z_i an i.i.d. vector hold in greater generality

From a technical viewpoint, analysis of the isotropic covariates model is straightforward and relies on standard random matrix theory results. However, we believe it provides useful insights.

In contrast, the results for general covariance and coefficients structure (Σ, β) is technically novel. We discuss related work in Section 1.3. Our results for the nonlinear model are also technically novel. In this setting, we derive a new asymptotic result on resolvents of certain block matrices, which may be of independent interest (see Lemma 3).

We next discuss the intuitions that emerge from our results as well as earlier literature.

Bias and variance. The shape of the asymptotic risk curve for min-norm least squares is, of course, controlled by its components: bias and variance. For fully specified models, the bias increases with γ in the overparametrized regime, which is intuitive. When $p > n$, the min-norm least squares estimate of β is constrained to lie the row space of X , the training feature matrix. This is a subspace of dimension n lying in a feature space of dimension p . Thus as p increases, so does the bias, since this row space accounts for less and less of the ambient p -dimensional feature space.

Meanwhile, we find that, in the overparametrized regime, the variance *decreases* with γ . This may seem counterintuitive at first, because it says, in a sense, that the min-norm least squares estimator becomes *more* regularized as p grows. However, this too can be explained intuitively, as follows. As p grows, the minimum ℓ_2 norm least squares solution—that is, the minimum ℓ_2 norm solution to the linear system $Xb = y$, for a training feature matrix X and response vector y —will generally have decreasing ℓ_2 norm. Why? Compare two such linear systems: in each, we are asking for the min-norm solution to a linear system with the same y , but in one instance we are given more columns in X , so we can generally decrease the components of b (by distributing them over more columns), and achieve a smaller ℓ_2 norm. This can in fact be formalized asymptotically; see Corollaries 1 and 3.

Double descent. Recently, Belkin et al. [8] pointed out a fascinating empirical trend where, for popular methods like neural networks and random forests, we can see a *second* bias-variance tradeoff in the out-of-sample prediction risk beyond the interpolation limit. The risk curve here resembles a traditional U-shape curve before the interpolation limit, and then descends again beyond the interpolation limit, which these authors call “double descent.” A closely related phenomenon was found earlier by Spigler et al. [63], who studied the “jamming transition” from underparametrized to overparametrized neural networks. Our results formally verify that this double descent phenomenon occurs even in the simple and fundamental case of least squares regression. The appearance of the second descent in the risk, past the interpolation boundary ($\gamma = 1$), is explained by the fact that the variance decreases as γ grows, as discussed above.

In the misspecified case, the variance still decreases with γ (for the same reasons), but interestingly, the bias can now also decrease with γ , provided γ is not too large (not too far past the interpolation boundary). The intuition here is that in a misspecified model, some part of the true regression function is always unaccounted for, and adding features generally improves our approximation capacity. As a consequence, the double descent phenomenon can be even more pronounced in the misspecified case (depending on the strength of the

approximation bias), in that the risk can attain its global minimum past the interpolation limit.

Finally, in the latent space model, we observe that the overall risk can be monotone decreasing in the overparametrized regime, and attain its global minimum for large overparametrization $\gamma \rightarrow \infty$ (after $p, n \rightarrow \infty$). In this case, we can write the design matrix as $X = ZW^T + U$, where U is noise, and Z is the $n \times d$ matrix of latent covariates. Equivalently, the i th column of X (the i th feature) takes the form $\tilde{x}_i = Zw_i + \tilde{u}_i$, where w_i is the i th column of W^T and $\tilde{u}_i$ is the i th column of U . Therefore, each new feature provides new information about the underlying low-dimensional latent variables Z . As p gets large, ridge regression with respect to the feature matrix X approximates increasingly well a ridge regression with respect to the latent variables Z .

Interpolation versus regularization. The min-norm least squares estimator can be seen as the limit of ridge regression as the tuning parameter tends to zero. A natural and important question is whether (or when) letting the regularization to 0 is optimal. Min-norm least squares is also the convergence point of gradient descent run on the least squares loss. Early-stopped gradient descent is known to be closely connected to ridge regularization; see, for example, Ali et al. [3], which proves a tight coupling between the two (see Section 1.3 for further related work). The question of whether letting the regularization vanish is optimal is closely related to the question of whether running gradient descent until convergence is optimal or early stopping provides some advantage.

Closely related questions have been investigated in the context of classification. For instance, it is common to run boosting until the training error is zero, and the boosting path is tied to ℓ_1 regularization [37, 58, 65]. It is empirically observed that, for noisy labels, early stopping (treating the number of boosting iterations as a tuning parameter) can be beneficial.

We would not expect the best-predicting ridge solution to be always at the end of its regularization path. Our results, comparing min-norm least squares to optimally-tuned ridge regression, show that (asymptotically) this is never the case, when β is incoherent with respect to the eigenvectors of Σ . This is for instance the case when $\Sigma = I_p$, or β is distributed according to a spherically symmetric prior. In contrast, [42] recently pointed out that—when β is aligned with the leading eigenvectors of Σ —min-norm regression can have optimal risk (i.e., the optimal regularization vanishes). We show that this is indeed the case in the latent space model mentioned above: this provides indeed an extremely simple example of a phenomenon that has been observed in the past for kernel methods [47]. Notice that the results we obtain for the latent space model are asymptotically sharper than the one of [47], in that they imply optimality up to subleading terms (not just up to constant factors).

In practice, of course, we would not have access to the optimal tuning parameter for ridge (optimal stopping for gradient descent), and we would rely on, for example, cross-validation (CV). Our theory shows that for ridge regression, CV tuning is asymptotically equivalent to optimal tuning (and we would expect the same results to carry over to gradient descent, but have not pursued this formally).

1.2. Connection to neural networks. As mentioned above, recent literature has established a direct connection between linear models and more complex models such as neural networks, in a certain training regime [4, 19, 24, 25, 39]. Here, we will briefly outline this connection, referring the reader to the literature for a more detailed exposition.

For the discussion in this section, it is convenient to consider a more general setting, in which we are given data (y_i, z_i) , $i \leq n$, $y_i \in \mathbb{R}$, $z_i \in \mathbb{R}^d$, which are i.i.d. from an arbitrary distribution $(y_i, z_i) \sim P_{y,z}$. Imagine training a neural network with parameters (weights) $\theta \in \mathbb{R}^p$, $f(\cdot; \theta) : \mathbb{R}^d \rightarrow \mathbb{R}$, $z \mapsto f(z; \theta)$. The specific form or architecture of the network is not important for our discussion. However, it is important to distinguish two conceptually different

functions. On the one hand, we have the true regression function $f_*(z_i) = \mathbb{E}\{y_i|z_i\}$; this is unknown to the statistician. For theoretical purposes, f_* can be assumed to belong to some function class, but for this section we will not specify this choice. On the other hand, we have a parametric model $f(z_i; \theta)$, which is determined by the specific network architecture. A specific network with the given architecture is determined by assigning the network weights $\theta \in \mathbb{R}^p$.

From the point of view of optimization, the central role is played by the parametric model $f(z; \theta)$. In modern machine learning, the number of parameters p is so large that—under certain training schemes— θ only changes by a small amount with respect to a random initialization $\theta_0 \in \mathbb{R}^p$. It thus makes sense to linearize the model around θ_0 . Supposing that the initialization is such that $f(z; \theta_0) \approx 0$, and letting $\theta = \theta_0 + \beta$, we can approximate the statistical model $z \mapsto f(z; \theta)$ by

$$(3) \quad z \mapsto \nabla_{\theta} f(z; \theta_0)^T \beta.$$

This model is still nonlinear in the input z , but is linear in the parameters β . In other words, if the linear approximation is accurate, learning reduces to computing feature vectors $x_i = \nabla_{\theta} f(z_i; \theta_0)$, $i = 1, \dots, n$, from the data, and then fitting a linear model in the x_i 's. Notice that these feature vectors have high dimension ($p > n$) since the network is overparametrized, and that the “featurization map” $z_i \mapsto \nabla f(z_i; \theta_0)$ is random because the initialization θ_0 is. Further, since $p > n$, many vectors β give rise to a model that interpolates the data.

The above scenario was made rigorous in a number of papers [4, 19, 24, 25, 39]. In particular, [19] shows—under some technical conditions—that the linearization (3) can be accurate if the model is overparametrized ($p > n$), and closed under scalings (if $f(\cdot)$ is encoded by a neural network, then $sf(\cdot)$ is also a neural network for any $s \in \mathbb{R}$). Under these conditions, there exists a scaling of the network's parameters such that gradient-based training converges to a model that can be approximated arbitrarily well by (3). Further, under the linearization (3), gradient descent converges to the interpolator that minimizes² the ℓ_2 norm $\|\beta\|_2$ (see Proposition 1 below).

What are the *statistical consequences* of these linearization results? In principle, we could consider $\{(y_i, z_i)\}_{i \leq n}$ to be i.i.d. samples with a certain population distribution $P_{y,z}$, and then study the behavior of minimum ℓ_2 norm interpolator of the form (3) under this data model. Denoting by $\phi(z) := \nabla_{\theta} f(z; \theta_0)$ the featurization map, this would amount to study

$$(4) \quad \hat{\beta} = \arg \min \{ \|\beta\|_2 \text{ subject to } x_i^T \beta = y_i, x_i = \phi(z_i) \forall i \leq n \}.$$

Even starting from a simple joint distribution $P_{y,z}$ for the data (y_i, z_i) , the resulting joint distribution for (y_i, x_i) (induced by the map $z_i \mapsto \phi(z_i) = \nabla_{\theta} f(z_i; \theta_0)$) can be very complicated.

From this point of view, the present paper establishes results for two types of featurization maps: in the linear model $\phi(z_i) = \Sigma^{1/2} z_i$, and in the nonlinear model $\phi(\bar{z}_i) = \varphi(W \bar{z}_i)$ where $W \in \mathbb{R}^{p \times d}$ and $\varphi: \mathbb{R} \rightarrow \mathbb{R}$ is applied componentwise. While both models are significantly simpler than the featurization map $\phi(z) = \nabla_{\theta} f(z; \theta_0)$ for a multilayer neural network, our results imply that certain universality phenomena hold in the proportional asymptotics $n, p, d \rightarrow \infty$, with $n \asymp p \asymp d$. Namely, under the assumption of z_i with independent coordinates in the linear model, and $\bar{z}_i \sim N(0, I_d)$ in the nonlinear model, we prove that:

1. If the activation function φ is “purely nonlinear” (in a sense to be made precise below), then the risk of the nonlinear model is asymptotically equal to the one of the linear model with $\Sigma = I_p$.

²Understanding the bias induced by gradient-based algorithms on fully nonlinear models is a broadly open problem, which has attracted considerable attention recently; see, for example, [34, 35].

2. For more general activations φ , we compute explicitly the asymptotics of the variance of the nonlinear model. This does not coincide with the variance in the isotropic model, but depends on φ only through the size of the linear component of φ (to be defined below). Once more, the details of the activation function do not matter.

After the present work appeared, the analysis of the nonlinear model was generalized in [49], which obtained the asymptotics of the risk for general activations, under a nonparametric model for the responses y_i . This required computing the bias term beyond purely nonlinear activations, and hence solving several technical challenges. The results of [49] confirmed that universality extends beyond purely nonlinear activations. For general activations, the nonlinear model with isotropic $\bar{z}_i \in \mathbb{R}^d$ is asymptotically equivalent to a linear model with anisotropic $x_i \in \mathbb{R}^p$ (analogous to the latent space model of Section 5.2). It is currently an open problem to which extent universality applies beyond the proportional regime.

Let us emphasize that universality is not expected to hold for any distribution of the data (y_i, z_i) , and for any function f . In particular, we not expect it to hold when z_i is low-dimensional. This is quite obvious from the proof of Theorem 8, and consistent with the findings of [56], which point at a qualitatively different behavior for interpolating methods in low dimension.

Finally, the correspondence outlined above only holds in a certain “lazy training” regime, in which network weights do not change much during training. More generally, in a neural network, the feature representation and the regression function or classifier are learned *simultaneously*. In terms of the first-order Taylor expansion (3) this means that θ_0 depends itself on the data, and hence the feature vectors $x_i = \nabla_{\theta} f(z_i; \theta_0)$ are not merely observed but trained. Learning the feature map could significantly change some aspects of the behavior of an interpolator. (See, for instance, Chapter 9 of Goodfellow et al. [33], and also Chizat and Bach [19], Zhang et al. [72], which emphasize the importance of learning the representation.)

1.3. Related work. The present work connects to and is motivated by the recent interest in interpolators in machine learning [8, 10, 11, 28, 47, 48]. Several authors have argued that minimum ℓ_2 norm least squares regression captures the basic behavior of deep neural networks, at least in early (lazy) training [4, 19, 24, 25, 39, 44, 73]. The connection between neural networks and kernel ridge regression arises when the number of hidden units diverges. The infinite width limit was also studied (beyond the linearized regime) in [18, 50, 59, 62].

Interpolation has a long history in signal processing, where it is a method of choice to reconstruct a subsampled signal. The overparametrized regime corresponds to the use of overcomplete dictionaries, and the minimum- ℓ_2 norm criterion was used for selecting a specific interpolator [21]. It was subsequently recognized that sparsity promoting interpolators provide better data representations [16].

Ridge regression with random designs has been studied in the past. Dicker [22] considers a model in which the covariates are isotropic Gaussian $x_i \sim N(0, I_p)$ and computes the asymptotic risk of ridge regression in the proportional asymptotics $p, n \rightarrow \infty$, with $p/n \rightarrow \gamma \in (0, \infty)$. Dobriban and Wager [23] generalize these results to $x_i = \Sigma^{1/2} z_i$, where z_i has independent entries with bounded 12th moment.

Recently, Advani and Saxe [2] study the effect of early stopping and ridge regularization in a model with isotropic Gaussian covariates $x_i \sim N(0, I_p)$, again focusing on the proportional asymptotics $p, n \rightarrow \infty$, with $p/n \rightarrow \gamma \in (0, \infty)$. They show that this simple model reproduces several phenomena observed in neural networks training. The same model is reconsidered in concurrent work by Belkin et al. [9], who obtain exact results for the expected risk of min-norm regression, relying on the jointly Gaussian distribution of (y_i, x_i) . We contribute to this line of work by extending the analysis to general covariance structures, non-Gaussian

covariates and to misspecified models. As we will see, these generalizations allow to produce examples for which the global minimum of the risk is achieved in the overparametrized regime $\gamma > 1$.

The importance of the relation between the coefficient vector β and the eigenvectors of Σ was emphasized by Kobak et al. [42] and Bartlett et al. [6]. These papers point out—under different asymptotic settings—that $\lambda = 0+$ (i.e., min-norm regression) can be optimal or nearly optimal. After a preprint of this paper appeared, Wu and Xu [69] and Richards et al. [57] generalized our earlier results to cover the case in which β is potentially aligned with Σ . We review in further detail these important generalizations in Section 4. We contribute to this line of work by obtaining nonasymptotic approximations for the risk, with explicit and nearly optimal error bounds. These hold under weaker assumptions on the geometry of (Σ, β) than the results of [57, 69].

High-dimensional regression under factor models for the covariates was recently studied by Bunea et al. [14], Bing et al. [12]. These models are related to the latent space model of Section 5.2 and present results complementary to ours.

For the nonlinear model, the random matrix theory literature is much sparser, and focuses on the related model of kernel random matrices, namely, symmetric matrices of the form $K_{ij} = \varphi(z_i^T z_j)$. El Karoui [26] studied the spectrum of such matrices in a regime in which φ can be approximated by a linear function (for $i \neq j$), and hence the spectrum converges to a rescaled Marchenko–Pastur law. This approximation does not hold for the regime of interest here, which was studied instead by Cheng and Singer [17] (who determined the limiting spectral distribution) and Fan and Montanari [27] (who characterized the extreme eigenvalues). The resulting eigenvalue distribution is the free convolution of a semicircle law and a Marchenko–Pastur law. In the current paper, we must consider asymmetric (rectangular) matrices $x_{ij} = \varphi(w_j^T z_i)$, whose singular value distribution was recently computed by Pennington and Worah [55], using the moment method. Unfortunately, the prediction variance depends on both the singular values and vectors of this matrix. In order to address this issue, we apply the leave-one out method of Cheng and Singer [17] to compute the asymptotics of the resolvent of a suitably extended matrix. We then extract the information of interest from this matrix. After appearance of a preprint of this paper, Mei and Montanari [49] extended the results presented here, to obtain a complete characterization of the risk for the nonlinear random features model.

Let us finally mention that the universality (or “invariance”) phenomenon is quite common in random matrix theory [64]. In the context of kernel inner product random matrices, it appears (somewhat implicitly) in [17] and (more explicitly) in [27]. After a first appearance of this manuscript, universality has been investigated in the context of neural networks in several papers [1, 29, 31, 38, 49, 52].

1.4. Outline. Section 2 provides important background. Sections 3–7 consider the linear model, focusing on isotropic features, correlated features, misspecified models, ridge regularization and cross-validation, respectively. Section 8 covers the nonlinear model case. Nearly all proofs are deferred until the Appendix.

2. Preliminaries. We describe our setup and gather a number of important preliminary results.

2.1. Data model and risk. Assume we observe training data $(x_i, y_i) \in \mathbb{R}^P \times \mathbb{R}$, $i = 1, \dots, n$ from the model of equations (1), (2). We collect the responses in a vector $y \in \mathbb{R}^n$, and the features in a matrix $X \in \mathbb{R}^{n \times P}$ (with rows $x_i \in \mathbb{R}^P$, $i = 1, \dots, n$).

Consider a test point $x_0 \sim P_x$, independent of the training data. For an estimator $\hat{\beta}$ (a function of the training data X, y), we define its out-of-sample prediction risk (or simply, risk) as

$$R_X(\hat{\beta}; \beta) = \mathbb{E}[(x_0^T \hat{\beta} - x_0^T \beta)^2 | X] = \mathbb{E}[\|\hat{\beta} - \beta\|_\Sigma^2 | X],$$

where $\|x\|_\Sigma^2 = x^T \Sigma x$. Note that our definition of risk is conditional on X (as emphasized by our notation R_X). Note also that we have the bias-variance decomposition

$$(5) \quad R_X(\hat{\beta}; \beta) = \underbrace{\|\mathbb{E}(\hat{\beta}|X) - \beta\|_\Sigma^2}_{B_X(\hat{\beta}; \beta)} + \underbrace{\text{Tr}[\text{Cov}(\hat{\beta}|X)\Sigma]}_{V_X(\hat{\beta}; \beta)}.$$

2.2. Ridgeless least squares. We consider the minimum ℓ_2 norm (min-norm) least squares regression estimator, of y on X , defined by

$$(6) \quad \hat{\beta} = \arg \min \{\|b\|_2 : b \text{ minimizes } \|y - Xb\|_2^2\}.$$

This can be equivalently written as $\hat{\beta} = (X^T X)^+ X^T y$, where $(X^T X)^+$ is the pseudoinverse of $X^T X$. An alternative name for (6) is the “ridgeless” least squares estimator, motivated by the fact that $\hat{\beta} = \lim_{\lambda \rightarrow 0^+} \hat{\beta}_\lambda$, where $\hat{\beta}_\lambda$ denotes the ridge regression estimator:

$$(7) \quad \hat{\beta}_\lambda = \arg \min_{b \in \mathbb{R}^p} \left\{ \frac{1}{n} \|y - Xb\|_2^2 + \lambda \|b\|_2^2 \right\},$$

or, equivalently, $\hat{\beta}_\lambda = (X^T X + n\lambda I)^{-1} X^T y$.

When X has full column rank the min-norm least squares estimator reduces to $\hat{\beta} = (X^T X)^{-1} X^T y$, the usual least squares estimator. When X has rank n , importantly, this estimator interpolates the training data: $y_i = x_i^T \hat{\beta}$, for $i = 1, \dots, n$.

Lastly, the following is a well-known fact that connects the min-norm least squares solution to gradient descent (as referenced in the [Introduction](#)).

PROPOSITION 1. *Initialize $\beta^{(0)} = 0$, and consider running gradient descent on the least squares loss, yielding iterates*

$$\beta^{(k)} = \beta^{(k-1)} + t X^T (y - X\beta^{(k-1)}), \quad k = 1, 2, 3, \dots,$$

where we take $0 < t \leq 1/\lambda_{\max}(X^T X)$ (and $\lambda_{\max}(X^T X)$ is the largest eigenvalue of $X^T X$). Then $\lim_{k \rightarrow \infty} \beta^{(k)} = \hat{\beta}$, the min-norm least squares solution in (6).

PROOF. The choice of step size guarantees that $\beta^{(k)}$ converges to a least squares solution as $k \rightarrow \infty$, call it $\tilde{\beta}$. Note that $\beta^{(k)}$, $k = 1, 2, 3, \dots$ all lie in the row space of X ; therefore, $\tilde{\beta}$ must also lie in the row space of X ; and the min-norm least squares solution $\hat{\beta}$ is the unique least squares solution with this property. $\square$

2.3. Bias and variance. We recall expressions for the bias and variance of the min-norm least squares estimator, which are standard.

LEMMA 1. *Under the model (1), (2), the min-norm least squares estimator (6) has bias and variance*

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

where $\hat{\Sigma} = X^T X/n$ is the (uncentered) sample covariance of X , and $\Pi = I - \hat{\Sigma}^+ \hat{\Sigma}$ is the projection onto the null space of X .

PROOF. As $\mathbb{E}(\hat{\beta}|X) = (X^T X)^+ X^T X \beta = \hat{\Sigma}^+ \hat{\Sigma} \beta$ and $\text{Cov}(\hat{\beta}|X) = \sigma^2 (X^T X)^+ X^T \times X (X^T X)^+ = \sigma^2 \hat{\Sigma}^+ / n$, the bias and variance expressions follow from plugging these into their respective definitions. $\square$

2.4. Underparametrized asymptotics. We consider an asymptotic setup where $n, p \rightarrow \infty$, in such a way that $p/n \rightarrow \gamma \in (0, \infty)$. Recall that when $\gamma < 1$, we call the problem underparametrized; when $\gamma > 1$, we call it overparametrized. Here, we recall the risk of the min-norm least squares estimator in the underparametrized case. The rest of this paper focuses on the overparametrized case.

The following is a known result in random matrix theory, and can be found in Chapter 6 of Serdobolskii [61]. It can also be found in the wireless communications literature; see Chapter 4 of Tulino and Verdu [67].

PROPOSITION 2. *Assume the model (1), (2), and assume $x \sim P_x$ is of the form $x = \Sigma^{1/2}z$, where z is a random vector with i.i.d. entries that have zero mean, unit variance and a finite 4th moment, and Σ is a (sequence of) deterministic positive definite matrix, such that $\lambda_{\min}(\Sigma) \geq c > 0$, for all n, p and a constant c (here $\lambda_{\min}(\Sigma)$ is the smallest eigenvalue of Σ). Then as $n, p \rightarrow \infty$, such that $p/n \rightarrow \gamma < 1$, the risk of the least squares estimator (6) satisfies, almost surely*

$$\lim_{n \rightarrow \infty} R_X(\hat{\beta}; \beta) = \sigma^2 \frac{\gamma}{1 - \gamma}.$$

As it can be seen from the last proposition, in the underparametrized case the risk is just variance. In contrast, in the overparametrized case, the bias $B_X(\hat{\beta}; \beta) = \beta^T \Pi \Sigma \Pi \beta$ is nonzero, because Π is. This will be the focus of the next sections.

3. Isotropic features. We begin by considering the simpler case in which $\Sigma = I$. In this case, the limiting bias is relatively straightforward to compute and depends on β only through $\|\beta\|_2^2$. In Section 4, we generalize our analysis and study the dependence of the prediction risk on the geometry of Σ and β .

3.1. Limiting bias. As mentioned above, in the isotropic case the risk depends β only on through $r^2 = \|\beta\|_2^2$. To give some intuition as to why this is true, consider the special case where X has i.i.d. entries from $N(0, 1)$. By rotational invariance, for any orthogonal $U \in \mathbb{R}^{p \times p}$, the distribution of X and XU is the same. Thus

$$\begin{aligned} B_X(\hat{\beta}; \beta) &= \beta^T (I - (X^T X)^+ X^T X) \beta \\ &\stackrel{d}{=} \beta^T (I - U^T (X^T X)^+ U U^T X^T X U) \beta \\ &= r^2 - (U\beta)^T (X^T X)^+ X^T X (U\beta). \end{aligned}$$

Choosing U so that $U\beta = r e_i$, the i th standard basis vector, then averaging over $i = 1, \dots, p$, yields

$$\mathbb{E} B_X(\hat{\beta}; \beta) = r^2 \mathbb{E}[1 - \text{Tr}((X^T X)^+ X^T X)/p] = r^2(1 - n/p).$$

It is possible to show that, $B_X(\hat{\beta}; \beta)$ concentrates around its expectation and, therefore, $B_X(\hat{\beta}; \beta) \rightarrow r^2(1 - 1/\gamma)$, almost surely. This is stated formally in the next section.

3.2. Limiting risk. As the next result shows, the independence of the risk on β is still true outside of the Gaussian case, provided the features are isotropic. The next result can be proved as a corollary of the more general Theorem 3 below. We give a simpler self-contained proof using a theorem of Rubio and Mestre [60] in Appendix A.4.2.

THEOREM 1. Assume the model (1), (2), where $x_i \sim P_x$ has independent entries with zero mean, unit variance. Further assume that either of these conditions hold for $x \sim P_x$: (i) the entries $(x_j)_{j \leq p}$ have uniformly bounded moments of all order $\mathbb{E}[|x_j|^k] \leq C_k$ for all k and some constants C_k ; (ii) entries $(x_j)_{j \leq p}$ are identically distributed and have finite moment of order $4 + \delta$, $\mathbb{E}\{|x_j|^{4+\delta}\} \leq C$, for some $C, \delta > 0$. Also assume that $\|\beta\|_2^2 = r^2$ for all n, p . Then for the min-norm least squares estimator $\hat{\beta}$ in (6), as $n, p \rightarrow \infty$, such that $p/n \rightarrow \gamma \in (1, \infty)$, it holds almost surely that

(8)
$$B_X(\hat{\beta}; \beta) \rightarrow r^2 \left(1 - \frac{1}{\gamma}\right),$$

(9)
$$V_X(\hat{\beta}; \beta) \rightarrow \sigma^2 \frac{1}{\gamma - 1}.$$

Hence, summarizing with Proposition 2, we have

(10)
$$R_X(\hat{\beta}; \beta) \rightarrow \begin{cases} \sigma^2 \frac{\gamma}{1 - \gamma} & \text{for } \gamma < 1, \\ r^2 \left(1 - \frac{1}{\gamma}\right) + \sigma^2 \frac{1}{\gamma - 1} & \text{for } \gamma > 1. \end{cases}$$

For $\gamma \in (0, 1)$, there is no bias, and the variance increases with γ . For $\gamma \in (1, \infty)$, the bias increases with γ , and the variance decreases with γ . Let $\text{SNR} = r^2/\sigma^2$. Observe that the risk of the null estimator $\tilde{\beta} = 0$ is r^2 , which we hence call the null risk. The following facts are immediate from the form of the risk curve in (10). See Figure 2 for an accompanying plot when SNR varies from 1 to 5.

- 1. On $(0, 1)$, the least squares risk $R(\gamma)$ is better than the null risk if and only if $\gamma < \frac{\text{SNR}}{\text{SNR}+1}$.
- 2. On $(1, \infty)$, when $\text{SNR} \leq 1$, the min-norm least squares risk $R(\gamma)$ is always worse than the null risk. Moreover, it is monotonically decreasing, and approaches the null risk (from above) as $\gamma \rightarrow \infty$.

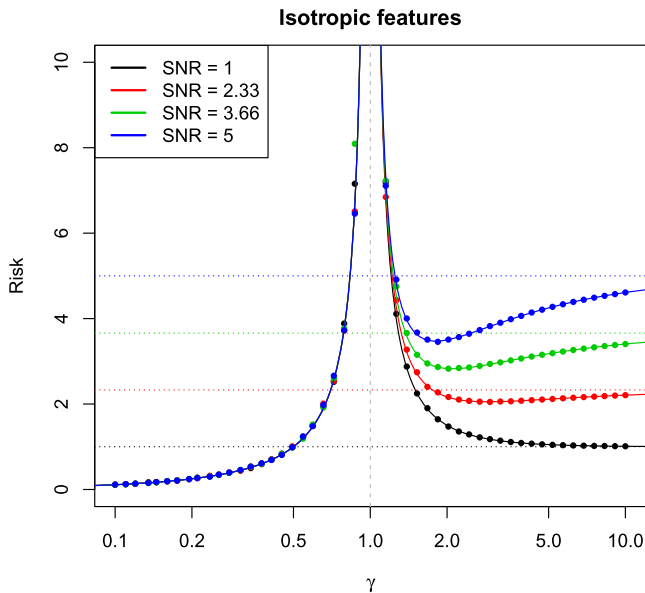


FIG. 2. Asymptotic risk curves in (10) for the min-norm least squares estimator, when r^2 varies from 1 to 5, and $\sigma^2 = 1$. For each value of r^2 , the null risk is marked as a dotted line, and the points denote finite-sample risks, with $n = 200$, $p = \lceil \gamma n \rceil$, across various values of γ , computed from features X having i.i.d. $N(0, 1)$ entries.

3. On $(1, \infty)$, when $\text{SNR} > 1$, the min-norm least squares risk $R(\gamma)$ beats the null risk if and only if $\gamma > \frac{\text{SNR}}{\text{SNR}-1}$. Further, it has a local minimum at $\gamma = \frac{\sqrt{\text{SNR}}}{\sqrt{\text{SNR}-1}}$, and approaches the null risk (from below) as $\gamma \rightarrow \infty$.

3.3. Limiting ℓ_2 norm. Calculation of the limiting ℓ_2 norm of the min-norm least squares estimator is quite similar to the study of the limiting risk in Theorem 1 and, therefore, we state the next result without proof.

COROLLARY 1. *Assume the conditions of Theorem 1. Then as $n, p \rightarrow \infty$, such that $p/n \rightarrow \gamma$, the squared ℓ_2 norm of the min-norm least squares estimator (6) satisfies, almost surely*

$$\mathbb{E}[\|\hat{\beta}\|_2^2 | X] \rightarrow \begin{cases} r^2 + \sigma^2 \frac{\gamma}{1-\gamma} & \text{for } \gamma < 1, \\ r^2 \frac{1}{\gamma} + \sigma^2 \frac{1}{\gamma-1} & \text{for } \gamma > 1. \end{cases}$$

We can see that the limiting norm, as a function of γ , has a somewhat similar profile to the limiting risk in (10): it is monotonically increasing for $\gamma \in (0, 1)$, diverges at the interpolation boundary and is monotonically decreasing for $(1, \infty)$. These findings confirm the intuition given in the [Introduction](#): as γ grows above the interpolation threshold, the minimum norm interpolator becomes increasingly simpler, in the sense of having smaller ℓ_2 norm.

4. Correlated features. We broaden the scope of our analysis from the last section, where we examined isotropic features. In this section, we take $x \sim P_x$ to be of the form $x = \Sigma^{1/2}z$, where z is a random vector with independent entries that have zero mean and unit variance, and Σ is arbitrary (but still deterministic and positive definite).

The risk of min-norm regression depends on the geometry of Σ and β . Denote by $\Sigma = \sum_{i=1}^p s_i v_i v_i^T$ the eigenvalue decomposition of Σ with $s_1 \geq s_2 \geq \dots \geq s_p \geq 0$. The geometry of the problem is captured by the sequence of eigenvalues $(s_1, \dots, s_p)$, and by the coefficients of β in the basis of eigenvectors $(\langle v_1, \beta \rangle, \dots, \langle v_p, \beta \rangle)$. We encode these via two probability distributions on $\mathbb{R}_{\geq 0}$:

$$(11) \quad \hat{H}_n(s) := \frac{1}{p} \sum_{i=1}^p 1_{\{s \geq s_i\}}, \quad \hat{G}_n(s) = \frac{1}{\|\beta\|_2^2} \sum_{i=1}^p \langle \beta, v_i \rangle^2 1_{\{s \geq s_i\}}.$$

We next state our assumptions about the data distribution: our results will be uniform with respect to the (large) constants $M, \{C_k\}_{k \geq 2}$ appearing in this assumption.

ASSUMPTION 1. The covariates vector $x \sim P_x$ is of the form $x = \Sigma^{1/2}z$, where defining $\hat{H}_n$ as per equation (11), we have:

- (a) The vector $z = (z_1, \dots, z_p)$ has independent (not necessarily identically distributed) entries with $\mathbb{E}\{z_i\} = 0$, $\mathbb{E}\{z_i^2\} = 1$, and $\mathbb{E}\{|z_i|^k\} \leq C_k < \infty$ for all $i \leq p, k \geq 2$.
- (b) $s_1 = \|\Sigma\|_{op} \leq M$, $\int s^{-1} d\hat{H}_n(s) < M$.
- (c) $|1 - (p/n)| \geq 1/M$, $1/M \leq p/n \leq M$.

Condition (a) bounds the tail probabilities on the covariates. Requiring finite moment of all orders is useful to get strong bounds on the deviations of the risk from its predicted value. As discussed below, bounds on the first few moments are sufficient if we are satisfied with weaker probability bounds.

Conditions (b) requires the eigenvalues of Σ to be bounded, and not to accumulate³ near 0. For the analysis of min-norm interpolation, we will add the additional assumption that the minimum eigenvalue of Σ is bounded away from zero. However, condition (b) is sufficient for the analysis of ridge regression in Section 6.

Finally, as our statements are nonasymptotic, we do not assume p/n to converge to a value. However, condition (c) requires p/n to be bounded and bounded away from the interpolation threshold $p/n = 1$.

4.1. Prediction risk.

DEFINITION 1 (Predicted bias and variance: min-norm regression). Let $\hat{H}_n$ be the empirical distribution of eigenvalues of Σ , and $\hat{G}_n$ the reweighted distribution as per equation (11). For $\gamma \in \mathbb{R}_{>0}$, define $c_0 = c_0(\gamma, \hat{H}_n) \in \mathbb{R}_{>0}$ to be the unique nonnegative solution of

$$(12) \quad 1 - \frac{1}{\gamma} = \int \frac{1}{1 + c_0 \gamma s} d\hat{H}_n(s).$$

We then define the predicted bias and variance by

$$(13) \quad \mathcal{B}(\hat{H}_n, \hat{G}_n, \gamma) := \|\beta\|_2^2 \left\{ 1 + \gamma c_0 \frac{\int \frac{s^2}{(1 + c_0 \gamma s)^2} d\hat{H}_n(s)}{\int \frac{s}{(1 + c_0 \gamma s)^2} d\hat{H}_n(s)} \right\} \cdot \int \frac{s}{(1 + c_0 \gamma s)^2} d\hat{G}_n(s),$$

$$(14) \quad \mathcal{V}(\hat{H}_n, \gamma) := \sigma^2 \gamma c_0 \frac{\int \frac{s^2}{(1 + c_0 \gamma s)^2} d\hat{H}_n(s)}{\int \frac{s}{(1 + c_0 \gamma s)^2} d\hat{H}_n(s)}.$$

Note that evaluating $\mathcal{B}(H, G, \gamma)$, $\mathcal{V}(H, \gamma)$ numerically is relatively straightforward, with the most complex part being the solution of equation (12). The next theorem establishes that—under suitable technical conditions—the functions $\mathcal{B}$, $\mathcal{V}$ characterize the test error. Similar theorems were proved in [57, 69], which generalized an earlier version of this manuscript to account for the geometry of (Σ, β) .

THEOREM 2. Assume the data model (1), (2) and that the covariates distribution satisfies Assumption 1. Further assume $s_p = \lambda_{\min}(\Sigma) > 1/M$. Define $\gamma = p/n$ and let $\hat{\beta}$ be the min-norm least squares estimator in equation (6).

Then for any constant $D > 0$ (arbitrarily large) there exist $C = C(M, D)$ such that, with probability at least $1 - Cn^{-D}$ the following hold:

$$(15) \quad R_X(\hat{\beta}; \beta) = B_X(\hat{\beta}; \beta) + V_X(\hat{\beta}; \beta),$$

$$(16) \quad |B_X(\hat{\beta}; \beta) - \mathcal{B}(\hat{H}_n, \hat{G}_n, \gamma)| \leq \frac{C \|\beta\|_2^2}{n^{1/7}},$$

$$(17) \quad |V_X(\hat{\beta}; \beta) - \mathcal{V}(\hat{H}_n, \gamma)| \leq \frac{C}{n^{1/7}},$$

where $\mathcal{B}$ and $\mathcal{V}$ are given in Definition 2, and the first identity is just the general bias-variance decomposition of equation (5).

The proof of this theorem is deferred to Section A.2.

³The latter assumption could have been further relaxed, by requiring only p_+ of the p eigenvalues to be non-vanishing and to satisfy the other conditions. This would require to redefine γ as p_+/n .

REMARK 1. The order of the error bound in equations (16), (17) is not optimal: a central limit theorem heuristics suggests the deterministic approximation to be accurate up to an error of order $n^{-1/2}$. Indeed, we are able to establish the optimal order in the case of ridge regression; see Theorem 5.

Let us emphasize that, while suboptimal, the $O(n^{-1/7})$ terms in equation (16), (17) are often negligible as compared to the leading terms $\mathcal{B}(\hat{H}_n, \hat{G}_n, \gamma)$, $\mathcal{V}(\hat{H}_n, \gamma)$. In particular, as stated in Theorem 3 below, whenever $p, n \rightarrow \infty$ with $p/n \rightarrow \gamma \in (0, \infty)$ and the two probability measures $\hat{H}_n, \hat{G}_n$ converge weakly to finite limits H, G , $\mathcal{B}(\hat{H}_n, \hat{G}_n, \gamma)$, $\mathcal{V}(\hat{H}_n, \gamma)$ remain bounded away from zero, and hence dominate the $O(n^{-1/7})$ errors.

Notice that this is in particular the case for isotropic features, and Theorem 1 (under the stronger moment assumption on the z_i 's) is recovered as a corollary of Theorem 2.

REMARK 2. Note that Theorem 2 establishes deterministic approximations for the bias and variance, that are valid at finite n, p . The overparametrization ratio $\gamma = p/n$ is a nonasymptotic quantity, and the error bounds are uniform, that is, depend only on the constant M . This is to be contrasted with the asymptotic setting of [57, 69]. Both of these papers assume a sequence of regression problems with $n, p \rightarrow \infty$, $p/n \rightarrow \gamma$, and obtain an asymptotically exact expression for the risk.

In order for the asymptotics to make sense, additional assumptions are required by [57, 69]. In [57], this is achieved by assuming β to be random with $\mathbb{E}[\beta\beta^T] = r^2\Phi(\Sigma)/d$ for a certain (deterministic) function $\Phi: \mathbb{R} \rightarrow \mathbb{R}$ (promoted to a function on matrices in the usual way). In addition, the empirical spectral distribution of Σ is assumed to converge. To state the assumptions in [69], recall that $(s_i)_{i \leq p}$ are the eigenvalues of Σ , and denote by $b_i = p \cdot (v_i^T \beta)^2$ the projection of β onto the eigenvectors of Σ . Then [69] assumes that the joint empirical distribution $p^{-1} \sum_{i=1}^p \delta_{s_i, b_i}$ converges weakly as $n, p \rightarrow \infty$.

Technically, [57, 69] apply asymptotic random matrix theory results, such as [43], while we have to take a longer detour to exploit nonasymptotic results established in [41]. We believe that the nonasymptotic approach provides more concrete and accurate statements.

Theorem 2 also implies asymptotic predictions under minimal assumptions. In particular, if the two probability measures $\hat{H}_n, \hat{G}_n$ converge weakly to probability measures H, G on $[0, \infty)$, then we obtain⁴ $B_X(\hat{\beta}; \beta) / \|\beta\|^2 \rightarrow \mathcal{B}(H, G, \gamma)$, $V_X(\hat{\beta}; \beta) \rightarrow \mathcal{V}(H, \gamma)$. Hence, the first part of the following asymptotic statement follows immediately from Theorem 2 by taking the limit $n, p \rightarrow \infty$ in equations (16), (17) (and using Borel–Cantelli to obtain almost sure convergence).

THEOREM 3. *Consider the setting of Theorem 2. Further assume $n, p \rightarrow \infty$, $p/n \rightarrow \gamma \in (0, \infty)$, $\hat{H}_n \Rightarrow H$, $\hat{G}_n \Rightarrow G$. Define $\mathcal{B}_1(H, G, \gamma)$ as in equation (13), with $\|\beta\|_2^2$ replaced by 1. Then, almost surely*

$$(18) \quad \frac{1}{\|\beta\|_2^2} B_X(\hat{\beta}; \beta) \rightarrow \mathcal{B}_1(H, G, \gamma), \quad V_X(\hat{\beta}; \beta) \rightarrow \mathcal{V}(H, \gamma).$$

with $\mathcal{B}_1(H, G, \gamma), \mathcal{V}(H, \gamma) > 0$ strictly.

The same conclusion holds if instead of Assumption 1(a), the coordinates of z , $(z_i)_{i \leq p}$ are i.i.d. and satisfy the conditions $\mathbb{E}z_i = 0$, $\mathbb{E}(z_i^2) = 1$, $\mathbb{E}(|z_i|^{4+\delta}) \leq C < \infty$.

⁴Indeed all the expressions in equations (12), (13), (14) are continuous in $\hat{H}_n, \hat{G}_n$ (with respect to the weak topology) since they are expectations of bounded continuous functions.

The last part of this theorem (under the weaker moment condition $\mathbb{E}(|z_i|^{4+\delta}) < C$) is proved via a truncation argument in Appendix A.1.4. For carrying out this argument, we make use of estimates on the norm of random matrices that are available only for the case of identically distributed entries $(z_i)_{i \leq p}$.

As pointed out above the condition $\widehat{H}_n \Rightarrow H$, $\widehat{G}_n \Rightarrow G$ (here $\Rightarrow$ denotes weak convergence) is strictly weaker than the condition assumed in [69] to establish asymptotic results. Further, we require weaker moment conditions.

In the next sections, we illustrate the role of the geometry of β , Σ by considering two models for which $\widehat{H}_n \Rightarrow H$, $\widehat{G}_n \Rightarrow G$ as $n, p \rightarrow \infty$. First, we consider the case $G = H$, which we refer to as “equidistributed”: the components of β are roughly equally distributed along the eigenvectors of Σ . In this case, there is no special relation between β and Σ .

As a further application, we consider a latent space model in which β is aligned with the top eigenvectors of Σ . This can be regarded as a misspecified model, and is therefore presented in Section 5.2 below.

4.2. Equidistributed coefficients. In this section, we assume $G = H$, $\|\beta\|_2 \rightarrow r$, and $p/n \rightarrow \gamma$. One way to generate β satisfying this condition is to draw it uniformly at random on the p -dimensional sphere of radius $\|\beta\|_2 = r$. In this case, the conditions of Theorem 2 (or Theorem 3), hold with $\widehat{G}_n \rightarrow G = H$.

COROLLARY 2. *Under the assumptions of Theorem 3, further assume $G = H$, $\|\beta\|^2 \rightarrow r^2$. Then for $n, p \rightarrow \infty$, with $p/n \rightarrow \gamma > 1$, almost surely*

$$(19) \quad B_X(\hat{\beta}; \beta) \rightarrow \mathcal{B}_{\text{equi}}(H, \gamma) := \frac{r^2}{c_0(H, \gamma)\gamma^2},$$

$$(20) \quad V_X(\hat{\beta}; \beta) \rightarrow \mathcal{V}_{\text{equi}}(H, \gamma) := \mathcal{V}(H, \gamma).$$

As a special case, we can revisit the isotropic case $\Sigma = I$, which results in $dH = \delta_1$. In this case, $c_0(H, \gamma) = \gamma(\gamma - 1)$ yielding immediately $\mathcal{B}_{\text{equi}}(H, \gamma) = 1 - \gamma^{-1}$ and $\mathcal{V}_{\text{equi}}(H, \gamma) = 1/(\gamma - 1)$.

4.3. Limiting ℓ_2 norm. Again, as in the isotropic case, analysis of the limiting ℓ_2 norm is similar to analysis of the risk in Theorem 2. We give the next result without proof, as it is an immediate generalization of previous results.

COROLLARY 3. *Under the assumptions of Theorem 3, further assume $\|\beta\|^2 \rightarrow r^2$, and let $c_0 = c_0(H, \gamma)$ be defined as there. Then as $n, p \rightarrow \infty$, such that $p/n \rightarrow \gamma$, the min-norm least squares estimator (6) satisfies, almost surely*

$$(21) \quad \|\hat{\beta}\|_2^2 \rightarrow \begin{cases} r^2 + \sigma^2 \frac{\gamma}{1-\gamma} \int \frac{1}{s} dH(s) & \text{for } \gamma < 1, \\ \int \frac{c_0 \gamma s}{1 + c_0 \gamma s} dG(s) + c_0 \gamma \sigma^2 & \text{for } \gamma > 1, \end{cases}$$

4.4. Benign overfitting. Theorem 2 (and its generalization to nonzero ridge regularization, Theorem 5) can be used to delineate regimes in which interpolation is statistically optimal or nearly optimal. “Statistical optimality” can be given different meanings in this context. Section 6.2 explores optimality in the context of the latent space model and shows that (in certain regimes) $R_X(\hat{\beta}_0; \beta) \leq (1 + o_n(1))R_X(\hat{\beta}_\lambda; \beta)$ for any $\lambda > 0$: min-norm interpolation is optimal up to subleading factors.

A different notion of optimality was explored (in concurrent work) in [6] and [66]. In these works, optimality is understood to hold up to constant multiplicative factors. A weaker notion is also considered whereby the interpolator is only required to be consistent. The term “benign overfitting” was proposed in [6] for such phenomena.

Theorem 2 can be used to establish upper bounds that are closely related to the ones of [6, 66] and in particular imply benign overfitting. As an example, the following bound was proven in joint work with Peter Bartlett and Alexander Rakhlin [7]. Recall that $s_1 \geq s_2 \geq \dots \geq s_p$ denote the eigenvalues of the covariance Σ in decreasing order, and we define the effective rank $r_k(\Sigma) = \sum_{i=k+1}^p \lambda_i / \lambda_{k+1}$. We also denote by $\beta_{\leq k}$ the projection of β onto the top k eigenvectors of Σ and $\beta_{>k} = \beta - \beta_{\leq k}$.

COROLLARY 4 ([7]). *Under the assumptions of Theorem 2, further assume that there exists an integer k and a constant $c_* > 0$ such that $r_k(\Sigma) \geq (1 + c_*)n$. Then there exists a constant $C = C(M, D)$ such that, with probability at least $1 - Cn^{-D}$,*

$$(22) \quad B_X(\hat{\beta}; \beta) \leq 4 \left(\frac{1}{n} \sum_{i=k+1}^p \lambda_i \right)^2 \|\beta_{\leq k}\|_{\Sigma^{-1}}^2 + \|\beta_{>k}\|_{\Sigma}^2 + Cn^{-1/7},$$

$$(23) \quad V_X(\hat{\beta}; \beta) \leq \frac{2k\sigma^2}{n} + \frac{4n\sigma^2}{c_*} \frac{\sum_{i=k+1}^p \lambda_i^2}{(\sum_{i=k+1}^p \lambda_i)^2} + Cn^{-1/7}.$$

This corollary is an immediate consequence of Theorem 2. It upper bounds the excess risk over an “ideal” (underparametrized) estimator that only fits the projection of β onto the top eigenvector of Σ . This excess risk will be small when β is well aligned to the top eigenvectors of Σ and the ratio $\sum_{i=k+1}^p \lambda_i^2 / (\sum_{i=k+1}^p \lambda_i)^2$ is small.

This result is analogous to the ones of [6, 66] although not precisely comparable. The upper bound on the variance term in [6] is sharp up to universal multiplicative constants. The upper bound on the bias in [66] depends on the (random) condition number of the component of the bias along less important directions. On the other hand, both of [6, 66] apply to cases in which $\Sigma^{-1/2}x$ does not have independent coordinates. Corollary 4 has a larger additive slack $Cn^{-1/7}$ (which can be improved to $Cn^{-1/2}$ for ridge regression), but a more precise prefactor.

5. Misspecified models.

5.1. Regression with respect to a subset of features. In this section, we consider a misspecified model, in which the regression function is still linear, but we observe only a subset of the features. Such a setting provides another potential motivation for interpolation: in many problems, we do not know the form of the regression function, and we generate features in order to improve our approximation capacity. Increasing the number of features past the point of interpolation (increasing γ past 1) can now decrease *both* bias and variance.

As such, the misspecified model setting also yields further interesting asymptotic comparisons between the $\gamma < 1$ and $\gamma > 1$ regimes. Recall the isotropic features model of Section 3.2: the risk function in (10) is globally minimized at $\gamma = 0$. This is a consequence of the fact that, in a well-specified linear model at $\gamma = 0$ there is no bias and no variance, and hence no risk. In a misspecified model, we will see that the story can be quite different, and the asymptotic risk can actually attain its *global* minimum on $(1, \infty)$.

5.1.1. *Data model and risk.* Consider, instead of (1), (2), a data model

$$(24) \quad ((x_i, w_i), \epsilon_i) \sim P_{x,w} \times P_\epsilon, \quad i = 1, \dots, n,$$

$$(25) \quad y_i = x_i^T \beta + w_i^T \theta + \epsilon_i, \quad i = 1, \dots, n,$$

where as before the random draws across $i = 1, \dots, n$ are independent. Here, we partition the features according to $(x_i, w_i) \in \mathbb{R}^{p+q}$, $i = 1, \dots, n$, where the joint distribution $P_{x,w}$ is such that $\mathbb{E}((x_i, w_i)) = 0$ and

$$\text{Cov}((x_i, w_i)) = \Sigma = \begin{bmatrix} \Sigma_x & \Sigma_{xw} \\ \Sigma_{xw}^T & \Sigma_w \end{bmatrix}.$$

We collect the features in a block matrix $[XW] \in \mathbb{R}^{n \times (p+q)}$ (which has rows $(x_i, w_i) \in \mathbb{R}^{p+q}$, $i = 1, \dots, n$). We presume that X is observed but W is unobserved, and focus on the min-norm least squares estimator exactly as before in (6), from the regression of y on X (not the full feature matrix $[XW]$).

Given a test point $(x_0, w_0) \sim P_{x,w}$, and an estimator $\hat{\beta}$ (fit using X , y only, and not W), we define its out-of-sample prediction risk as

$$R_X(\hat{\beta}; \beta, \theta) = \mathbb{E}[(x_0^T \hat{\beta} - \mathbb{E}(y_0|x_0, w_0))^2 | X] = \mathbb{E}[(x_0^T \hat{\beta} - x_0^T \beta - w_0^T \theta)^2 | X].$$

Note that this definition is conditional on X , and we are integrating over the randomness not only in ϵ (the training errors), but in the unobserved features W , as well. The next lemma decomposes this notion of risk in a useful way.

LEMMA 2. *Under the misspecified model (24), (25), for any estimator $\hat{\beta}$, we have*

$$R_X(\hat{\beta}; \beta, \theta) = \underbrace{\mathbb{E}[(x_0^T \hat{\beta} - \mathbb{E}(y_0|x_0))^2 | X]}_{R_X^*(\hat{\beta}; \beta, \theta)} + \underbrace{\mathbb{E}[(\mathbb{E}(y_0|x_0) - \mathbb{E}(y_0|x_0, w_0))^2]}_{M(\beta, \theta)}.$$

PROOF. Simply add an subtract $\mathbb{E}(y_0|x_0)$ inside the square in the definition of $R_X(\hat{\beta}; \beta, \theta)$, then expand, and note that the cross term vanishes because $\mathbb{E}[(\mathbb{E}(y_0|x_0) - \mathbb{E}(y_0|x_0, w_0))|x_0] = 0$. $\square$

The first term $R_X^*(\hat{\beta}; \beta, \theta)$ in the decomposition in Lemma 2 is precisely the risk that we studied previously in the well-specified case, except that the response distribution has changed (due to the presence of the middle term in (25)). We call the second term $M(\beta, \theta)$ in Lemma 2 the *misspecification bias*.

REMARK 3. If (x_i, w_i) are jointly Gaussian, then the above expressions simplify and Theorem 2 can be used to characterize the risk $R_X(\hat{\beta}; \beta, \theta)$. In particular, the conditional distribution of w given x is $P_{w|x} = N(\Sigma_{wx} \Sigma_x^{-1} x, \Sigma_{w|x})$ where $\Sigma_{wx} = \Sigma_{xw}^T$, and $\Sigma_{w|x} = \Sigma_w - \Sigma_{wx} \Sigma_x^{-1} \Sigma_{xw}$. Further, $y = \tilde{\beta}^T x + \tilde{\epsilon}$, where $\tilde{\beta} = \beta + \Sigma_x^{-1} \Sigma_{xw} \theta$ and $\tilde{\epsilon} \sim N(0, \tilde{\sigma}^2)$, $\tilde{\sigma}^2 = \sigma^2 + \theta^T \Sigma_{w|x} \theta$. It is then easy to show that the misspecification bias is $M(\beta, \theta) = \theta^T \Sigma_{w|x} \theta$ and the term $R_X^*(\hat{\beta}; \beta, \theta)$ can be approximated using Theorem 2.

In order to discuss some qualitative features, we focus on the simplest possible model by assuming independent covariates.

5.1.2. *Isotropic features.* Here, we make the additional simplifying assumption that $(x, w) \sim P_{x,w}$ has i.i.d. entries with unit variance, which implies that $\Sigma = I$. (The case of independent features but general covariances Σ_x, Σ_w is similar, and we omit the details.) Therefore, we may write the response distribution in (25) as

$$y_i = x_i^T \beta + \delta_i, \quad i = 1, \dots, n,$$

where δ_i is independent of x_i , having mean zero and variance $\sigma^2 + \|\theta\|_2^2$, for $i = 1, \dots, n$. Denote the total signal by $r^2 = \|\beta\|_2^2 + \|\theta\|_2^2$, and the fraction of the signal captured by the observed features by $\kappa = \|\beta\|_2^2 / r^2$. Then $R_X^*(\hat{\beta}; \beta, \theta)$ behaves exactly as we computed previously, for isotropic features in the well-specified setting (Theorem 2 for $\gamma < 1$, and Theorem 1 for $\gamma > 1$), after we make the substitutions:

$$(26) \quad r^2 \mapsto r^2 \kappa \quad \text{and} \quad \sigma^2 \mapsto \sigma^2 + r^2(1 - \kappa).$$

Furthermore, we can easily calculate the misspecification bias:

$$M(\beta, \theta) = \mathbb{E}(w_0^T \theta)^2 = r^2(1 - \kappa).$$

Putting these results together leads to the next conclusion.

THEOREM 4. *Assume the misspecified model (24), (25) and assume $(x, w) \sim P_{x,w}$ has i.i.d. entries with zero mean, unit variance and a finite moment of order $4 + \delta$, for some $\delta > 0$. Also assume that $\|\beta\|_2^2 + \|\theta\|_2^2 = r^2$ and $\|\beta\|_2^2 / r^2 = \kappa$ for all n, p . Then for the min-norm least squares estimator $\hat{\beta}$ in (6), as $n, p \rightarrow \infty$, with $p/n \rightarrow \gamma$, it holds almost surely that*

$$R_X(\hat{\beta}; \beta, \theta) \rightarrow \begin{cases} r^2(1 - \kappa) + (r^2(1 - \kappa) + \sigma^2) \frac{\gamma}{1 - \gamma} & \text{for } \gamma < 1, \\ r^2(1 - \kappa) + r^2 \kappa \left(1 - \frac{1}{\gamma}\right) + (r^2(1 - \kappa) + \sigma^2) \frac{1}{\gamma - 1} & \text{for } \gamma > 1. \end{cases}$$

We remark that, in the independence setting considered in Theorem 4, the dimension q of the unobserved feature space does not play any role: we may equally well take $q = \infty$ for all n, p (i.e., infinitely many unobserved features).

The components of the limiting risk from Theorem 4 are intuitive and can be interpreted as follows. The first term $r^2(1 - \kappa)$ is the misspecification bias (irreducible). The second term, which we deem as 0 for $\gamma < 1$ and $r^2 \kappa(1 - 1/\gamma)$ for $\gamma > 1$, is the bias. The third term, $r^2(1 - \kappa)\gamma/(1 - \gamma)$ for $\gamma < 1$ and $r^2(1 - \kappa)/(\gamma - 1)$ for $\gamma > 1$, is what we call the *misspecification bias*: the inflation in risk due to unobserved features, when we take $\mathbb{E}(y_0|x_0)$ to be the target of estimation. The last term, $\sigma^2\gamma/(1 - \gamma)$ for $\gamma < 1$ and $\sigma^2/(\gamma - 1)$ for $\gamma > 1$, is the variance itself.

5.1.3. *Polynomial approximation bias.* Since adding features should generally improve our approximation capacity, it is reasonable to model $\kappa = \kappa(\gamma)$ as an increasing function of γ . To get an idea of the possible shapes taken by the asymptotic risk curve from Theorem 4, we consider the example of a *polynomial decay* for the approximation bias,

$$(27) \quad 1 - \kappa(\gamma) = (1 + \gamma)^{-a},$$

for some $a > 0$. In this case, the limiting risk in the isotropic setting, from Theorem 4, becomes

$$(28) \quad R_a(\gamma) = \begin{cases} r^2(1 + \gamma)^{-a} + (r^2(1 + \gamma)^{-a} + \sigma^2) \frac{\gamma}{1 - \gamma} & \text{for } \gamma < 1, \\ r^2(1 + \gamma)^{-a} + r^2(1 - (1 + \gamma)^{-a}) \left(1 - \frac{1}{\gamma}\right) \\ \quad + (r^2(1 + \gamma)^{-a} + \sigma^2) \frac{1}{\gamma - 1} & \text{for } \gamma > 1. \end{cases}$$

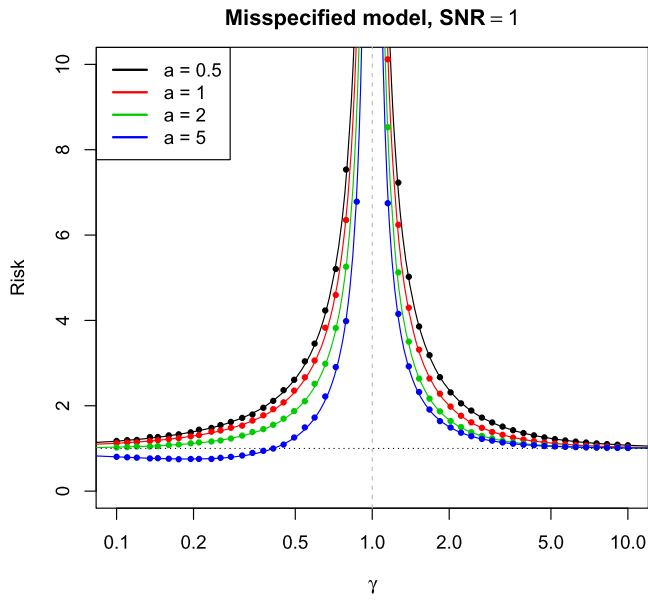


FIG. 3. Asymptotic risk curves in (28) for the min-norm least squares estimator in the misspecified case, when the approximation bias has polynomial decay as in (27), as a varies from 0.5 to 5. Here, $r^2 = 1$ and $\sigma^2 = 1$, so $\text{SNR} = 1$. The null risk $r^2 = 1$ is marked as a dotted black line. The points denote finite-sample risks, with $n = 200$, $p = \lceil \gamma n \rceil$, across various values of γ , computed from features X having i.i.d. $N(0, 1)$ entries.

We next summarize some interesting features of these risk curves, and Figures 3 and 4 give accompanying plots for $\text{SNR} = 1$ and 5, respectively. Recall that the null risk is r^2 , which comes from predicting with the null estimator $\tilde{\beta} = 0$.

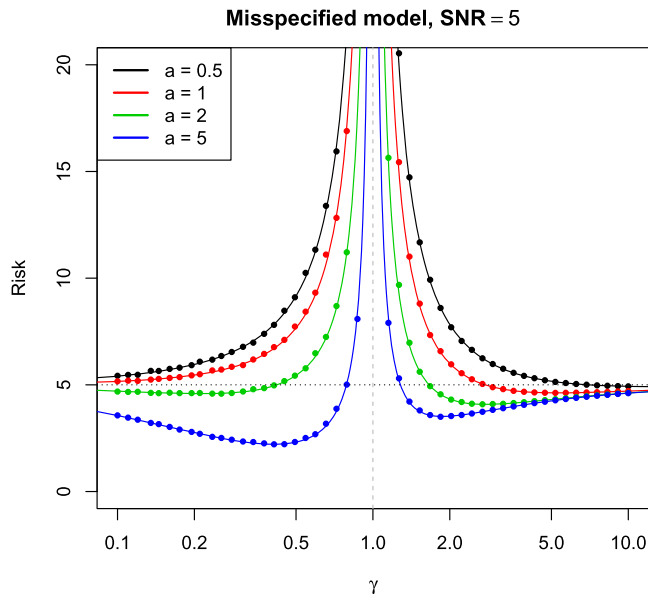


FIG. 4. Asymptotic risk curves in (28) for the min-norm least squares estimator in the misspecified case, when the approximation bias has polynomial decay as in (27), as a varies from 0.5 to 5. Here, $r^2 = 5$ and $\sigma^2 = 1$, so $\text{SNR} = 5$. The null risk $r^2 = 5$ is marked as a dotted black line. The points are again finite-sample risks, with $n = 200$, $p = \lceil \gamma n \rceil$, across various values of γ .

1. On $(0, 1)$, the least squares risk $R_a(\gamma)$ can only be better than the null risk if $a > 1 + \frac{1}{\text{SNR}}$. Further, in this case, we have $R_a(\gamma) < r^2$ if and only if $\gamma < \gamma_0$, where γ_0 is the unique zero of the function

$$(1+x)^{-a} + \left(1 + \frac{1}{\text{SNR}}\right)x - 1$$

that lies in $(0, \frac{\text{SNR}}{\text{SNR}+1})$. Finally, on $(\frac{\text{SNR}}{\text{SNR}+1}, 1)$, the least squares risk $R_a(\gamma)$ is always worse than the null risk, regardless of $a > 0$, and it is monotonically increasing.

2. On $(1, \infty)$, when $\text{SNR} \leq 1$, the min-norm least squares risk $R_a(\gamma)$ is always worse than the null risk. Moreover, it is monotonically decreasing, and approaches the null risk (from above) as $\gamma \rightarrow \infty$.

3. On $(1, \infty)$, when $\text{SNR} > 1$, the min-norm least squares risk $R_a(\gamma)$ can be better than the null risk for any $a > 0$, and in particular we have $R_a(\gamma) < r^2$ if and only if $\gamma < \gamma_0$, where γ_0 is the unique zero of the function

$$(1+x)^{-a}(2x-1) + 1 - \left(1 - \frac{1}{\text{SNR}}\right)x$$

lying in $(\frac{\text{SNR}}{\text{SNR}-1}, \infty)$. Indeed, on $(1, \frac{\text{SNR}}{\text{SNR}-1})$, the min-norm least squares risk $R_a(\gamma)$ is always worse than the null risk (regardless of $a > 0$), and it is monotonically decreasing.

4. When $\text{SNR} > 1$, for small enough $a > 0$, the global minimum of the min-norm least squares risk $R_a(\gamma)$ occurs after $\gamma = 1$. A sufficient but not necessary condition is $a \leq 1 + \frac{1}{\text{SNR}}$ (because, due to points 1 and 3 above, we see that in this case $R_a(\gamma)$ is always worse than null risk for $\gamma < 1$, but will be better than the null risk at some $\gamma > 1$).

5.2. Latent space model. We next consider an example in which β is aligned with the top eigenvectors of Σ . To motivate this example, assume that the responses y_i are linear in the latent features vectors $z_i \in \mathbb{R}^d$. We do not observe this latent vector, but rather observe $p \geq d$ covariates $x_i := (x_{i1}, \dots, x_{ip})$ that are also linear in the latent vector z_i :

$$(29) \quad x_i = \theta^T z_i + \xi_i, \quad x_{ij} = w_j^T z_i + u_{ij}.$$

Here, $(\xi_i)_{i \leq n}$, $(u_{ij})_{i \leq n, j \leq p}$ are noise variables that are mutually independent, and independent of z_i , with $\xi_i \sim N(0, \sigma_\xi^2)$, $u_{ij} \sim N(0, 1)$. The features matrix takes the form $X = ZW^T + U$ and, therefore, for $p > n$, min-norm regression amounts to

$$(30) \quad \hat{\beta} = \arg \min \{ \|b\|_2 : ZW^T b + Ub = y \}.$$

Apart from its intrinsic interest, this latent-space model is directly connected to nonlinear random features models, as the ones studied in Section 8. Indeed, in nonlinear random features models we have $x_{ij} = \varphi(w_j^T z_i)$. We can decompose this as $x_{ij} = a_0 + a_1 w_j^T z_i + \tilde{\varphi}(w_j^T z_i)$, where $\tilde{\varphi}$ is such that $\tilde{\varphi}(w_j^T z_i)$ has zero mean and is uncorrelated with $w_j^T z_i$, conditional on w_j . Equation (29) then corresponds to replacing the uncorrelated random variable $\tilde{\varphi}(w_j^T z_i)$ by the independent Gaussians u_{ij} . This connection was discussed in [49, 52], after a first appearance of the present paper. Recent studies of high-dimensional linear discriminant analysis [15, 40, 53] share important elements of the model studied here: anisotropic covariance and signal aligned with its top eigenvectors.

We consider a variant of model (29), which is a special case of the model studied in Section 4. Namely we assume $x_i = \Sigma^{1/2} \tilde{z}_i$, $y_i = \beta^T x_i + \varepsilon_i$, where $\tilde{z}_i$ is a vector with independent coordinates satisfying Assumption 1, and we set

$$(31) \quad \Sigma = I_p + WW^T, \quad \beta = W(I + W^T W)^{-1} \theta,$$

$$(32) \quad \mathbb{E}(\varepsilon_i) = 0, \quad \mathbb{E}(\varepsilon_i^2) = \sigma^2, \quad \sigma^2 = \sigma_\xi^2 + \theta^T (I + W^T W)^{-1} \theta.$$

Here, $W \in \mathbb{R}^{p \times d}$ is the matrix with rows $(w_i)_{i \leq p}$. In what follows, $r_\theta^2 := \|\theta\|_2^2$, $\psi = d/p$. As anticipated, the coefficients vector is aligned with the top eigenspace of Σ (the span of the columns of W).

The connection between the last formulation and the model of equation (29) is easy to see if the latent vector $z_i \sim N(0, I_d)$. In this case, the two models coincide because $(y_i, x_i) \in \mathbb{R}^{p+1}$ is a centered Gaussian vector with the same covariance structure.

In order to simplify our calculations, we assume all the nonzero singular values of W to be equal, whence $W^T W = (p\mu/d)I_d$, for $\mu > 0$ a constant. The factor p/d is justified by the remark that the average norm of the vectors w_j is given by

$$\frac{1}{p} \sum_{j=1}^p \|w_j\|_2^2 = \frac{1}{p} \text{tr}(W^T W) = \mu.$$

Hence, μ is the signal-to-noise ratio in the features $x_{ij} = w_j^T z_i + u_{ij}$, and keeping μ constant corresponds to keeping this signal-to-noise ratio constant. This is also motivated by the nonlinear random features model $x_{ij} = \varphi(w_j^T z_i)$ (if we identify the nonlinear component with the noise); see Section 8 and [49].

Note that with this setting, the eigenvalues of Σ are

$$s_1 = s_2 = \dots = s_d = 1 + \mu\psi^{-1} > s_{d+1} = s_2 = \dots = s_p = 1.$$

If $p, d, n \rightarrow \infty$, with $p/n \rightarrow \gamma$, $d/p \rightarrow \psi$, then this model satisfies the assumptions of Theorem 2, with

$$(33) \quad H(s) = (1 - \psi)\mathbf{1}(s \geq 1) + \psi\mathbf{1}(s \geq 1 + \psi^{-1}),$$

$$(34) \quad G(s) = \mathbf{1}(s \geq 1 + \psi^{-1}), \|\beta\|_2^2 = \frac{\mu\psi^{-1}r_\theta^2}{(1 + \mu\psi^{-1})^2}$$

Using Theorem 2, we get the following explicit expressions.

COROLLARY 5. *Consider the latent space model described above, namely $x_i = \Sigma^{1/2}\tilde{z}_i$, $y_i = \beta^T x_i + \varepsilon_i$, where $\tilde{z}_i$ is a vector with independent coordinates satisfying Assumption 1. Further assume equations (31), (32) and $d/p \rightarrow \psi \in (1, \infty)$, $p/n \rightarrow \gamma \in (1, \infty)$ to hold. (The case $\gamma \in (0, 1)$ being covered by Proposition 2.)*

Then almost surely

$$(35) \quad R_X(\hat{\beta}; \beta) \rightarrow \mathcal{B}_{\text{lat}}(\psi, \gamma) + \mathcal{V}_{\text{lat}}(\psi, \gamma),$$

$$(36) \quad \mathcal{B}_{\text{lat}}(\psi, \gamma) := \left\{ 1 + \gamma c_0 \frac{\mathcal{E}_1(\psi, \gamma)}{\mathcal{E}_2(\psi, \gamma)} \right\} \cdot \frac{\mu\psi^{-1}r_\theta^2}{(1 + \mu\psi^{-1})(1 + c_0\gamma(1 + \mu\psi^{-1}))^2},$$

$$(37) \quad \mathcal{V}_{\text{lat}}(\psi, \gamma) := \sigma^2 \gamma c_0 \frac{\mathcal{E}_1(\psi, \gamma)}{\mathcal{E}_2(\psi, \gamma)},$$

$$(38) \quad \mathcal{E}_1(\psi, \gamma) := \frac{1 - \psi}{(1 + c_0\gamma)^2} + \frac{\psi(1 + \mu\psi^{-1})^2}{(1 + c_0(1 + \mu\psi^{-1})\gamma)^2},$$

$$(39) \quad \mathcal{E}_2(\psi, \gamma) := \frac{1 - \psi}{(1 + c_0\gamma)^2} + \frac{\psi(1 + \mu\psi^{-1})}{(1 + c_0(1 + \mu\psi^{-1})\gamma)^2}.$$

where $\sigma^2 = \sigma_\varepsilon^2 + r_\theta^2/(1 + \mu\psi^{-1})$, and $c_0 = c_0(\psi, \gamma) \geq 0$ is the unique nonnegative solution of the following second-order equation:

$$(40) \quad 1 - \frac{1}{\gamma} = \frac{1 - \psi}{1 + c_0\gamma} + \frac{\psi}{1 + c_0(1 + \mu\psi^{-1})\gamma}.$$

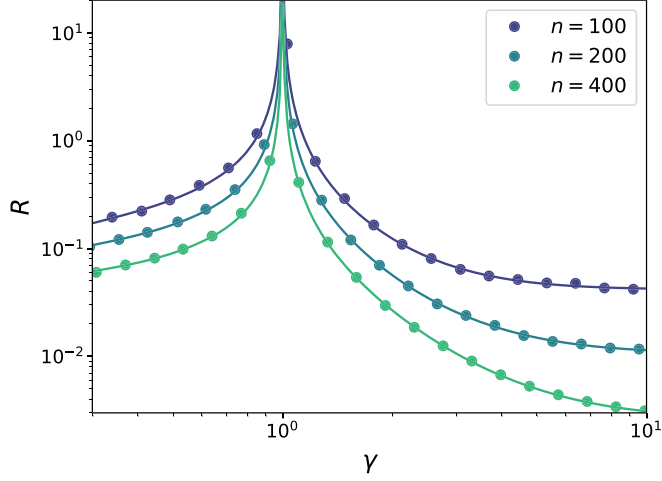


FIG. 5. Latent space model of Section 5.2: test error $R_X(\hat{\beta}; \beta)$ of minimum norm regression as a function of the overparametrization ratio γ . Here, $d = 20$, $r = 1$, $\mu = 1$, $\sigma_\xi = 0$ and n varies across various curves. Symbols are averages over 100 realizations; continuous lines report the analytical prediction of Corollary 5.

REMARK 4. The proof of Theorem 2 holds almost unchanged for the case in which $x_i = \Sigma^{1/2} \tilde{z}_i$ with $\Sigma^{1/2}$ a nonsymmetric square root of Σ and $\tilde{z}_i$ a vector with independent entries satisfying Assumption 2. This case includes the general model equation (29) for z_i with independent entries as a special case. It is sufficient to set $\tilde{z}_i = (z_i, u_i)$ and $\Sigma^{1/2} = (W, I_p)$.

Figures 5 and 6 illustrate this corollary by comparing analytical predictions to numerical simulations. We observe that the prediction risk is monotone decreasing in the overparametrization ratio for $\gamma > 1$, and reaches its global minimum asymptotically as $\gamma \rightarrow \infty$ (after $p, n, d \rightarrow \infty$). To understand why this happens, notice that each feature vector x_i can be viewed as a noisy measurement of the latent covariates z_i . If the noise u_{ij} was absent, then performing min-norm regression with respect to $(x_i)_{i \leq n}$ would be equivalent to min-norm regression with respect to $(z_i)_{i \leq n}$. To see this, consider again equation (30). If we

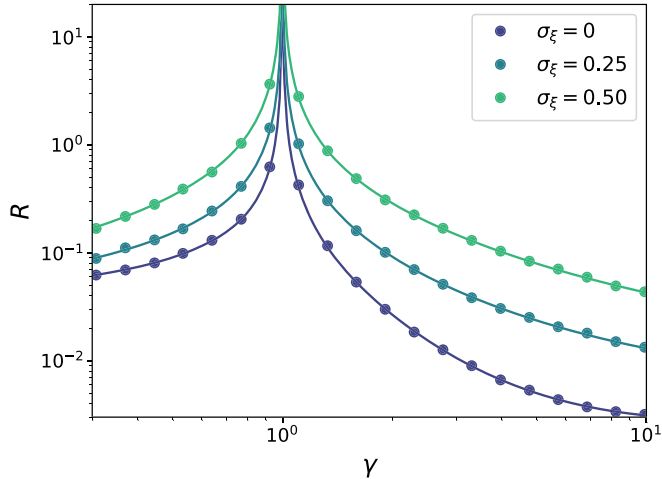


FIG. 6. Latent space model of Section 5.2: test error $R_X(\hat{\beta}; \beta)$ of minimum norm regression as a function of the overparametrization ratio γ . Here, $d = 20$, $r = 1$, $\mu = 1$, $n = 400$ and the noise variance σ_ξ varies across different curves. Symbols are averages over 100 realizations; continuous lines report the analytical prediction of Corollary 5.

drop the noise U , we are minimizing $\|b\|_2$ subject to $Z(W^T b) = 0$, and the regression function is $\hat{f}(z) = x^T \hat{\beta} = z^T (W^T \hat{\beta})$. Since W is orthogonal, this is equivalent to computing $\hat{\theta} = \arg \min\{\|t\|_2 : \text{subject to } Zt = y\}$, with $y = Z\theta + \xi$. In other words, we are back to the underparametrized model.

In presence of noise u_{ij} , the latent features cannot be estimated exactly. However, as p gets larger, the noise is effectively “averaged out” and we approach the idealized situation in which the z_i ’s are observed.

All of the simulations in Figures 5, 6 are carried out with $\mu = 1$. In Appendix A.3, we explore the dependence on μ , and show that the generalization curves are insensitive over a broad range of choices of this parameter.

6. Ridge regularization. We generalize the formulas of Section 4 to nonvanishing ridge regularization. We work under the same assumptions of that section. In particular, recall that $\hat{H}_n(s) = p^{-1} \sum_{i=1}^p 1_{\{s \geq s_i\}}$ is the empirical distribution of the eigenvalues of Σ , and $\hat{G}_n(s) = \sum_{i=1}^p \langle \beta, v_i \rangle^2 1_{\{s \geq s_i\}} / \|\beta\|^2$ the same empirical distribution, reweighted by the projection of β onto the eigenvectors v_i of the covariance Σ . (Recall the eigenvalue decomposition $\Sigma = \sum_{i=1}^p s_i v_i v_i^T$.)

DEFINITION 2 (Predicted bias and variance: ridge regression). For $\gamma \in \mathbb{R}_{>0}$, and $z \in \mathbb{C}_+$ (the set of complex numbers with $\text{Im}(z) > 0$), define $m_n(z) = m(z; \hat{H}_n, \gamma)$ as the unique solution of

$$(41) \quad m_n(z) = \int \frac{1}{s[1 - \gamma - \gamma z m_n(z)] - z} d\hat{H}_n(s).$$

Further define $m_{n,1}(z) = m_{n,1}(z; \hat{H}_n, \gamma)$ via

$$(42) \quad m_{n,1}(z) := \frac{\int \frac{s^2[1 - \gamma - \gamma z m_n(z)]}{[s[1 - \gamma - \gamma z m_n(z)] - z]^2} d\hat{H}_n(s)}{1 - \gamma \int \frac{zs}{[s[1 - \gamma - \gamma z m_n(z)] - z]^2} d\hat{H}_n(s)}.$$

These definitions are extended analytically to $\text{Im}(z) = 0$ whenever possible. We then define the predicted bias and variance by

$$(43) \quad \mathcal{B}(\lambda; \hat{H}_n, \hat{G}_n, \gamma) := \lambda^2 \|\beta\|^2 (1 + \gamma m_{n,1}(-\lambda)) \times \int \frac{s}{[\lambda + (1 - \gamma + \gamma \lambda m_n(-\lambda))s]^2} d\hat{G}_n(s),$$

$$(44) \quad \mathcal{V}(\lambda; \hat{H}_n, \gamma) := \sigma^2 \gamma \int \frac{s^2(1 - \gamma + \gamma \lambda^2 m'_n(-\lambda))}{[\lambda + s(1 - \gamma + \gamma \lambda m_n(-\lambda))]^2} d\hat{H}_n(s).$$

We next state our deterministic approximation of the risk.

THEOREM 5. Let $M^{-1} \leq p/n \leq M$, and Assumption 1 hold. Further assume $\lambda \vee s_{\min}(\Sigma) > 1/M$ and $n^{-2/3+1/M} < \lambda < M$. Let $\hat{\beta}_\lambda$ be the ridge estimator of equation (7).

Then for any constants $D > 0$ (arbitrarily large) and $\varepsilon > 0$ (arbitrarily small), there exist $C = C(M, D)$ such that, with probability at least $1 - Cn^{-D}$ the following hold:

$$(45) \quad R_X(\hat{\beta}_\lambda; \beta) = B_X(\hat{\beta}_\lambda; \beta) + V_X(\hat{\beta}_\lambda; \beta),$$

$$(46) \quad |B_X(\hat{\beta}_\lambda; \beta) - \mathcal{B}(\lambda; \hat{H}_n, \hat{G}_n, \gamma)| \leq \frac{C \|\beta\|_2^2}{\lambda n^{(1-\varepsilon)/2}},$$

$$(47) \quad |V_X(\hat{\beta}_\lambda; \beta) - \mathcal{V}(\lambda; \hat{H}_n, \gamma)| \leq \frac{C}{\lambda^2 n^{(1-\varepsilon)/2}},$$

where $\mathcal{B}$ and $\mathcal{V}$ are given in Definition 2, and the first identity is just the general bias-variance decomposition of equation (5).

The proof of this theorem is deferred to Appendix A.1. As for Theorem 2, similar results were proved in [57, 69], subsequently to a first version of this manuscript that only focused on random β . The same comparison of Remark 2 applies here.

In particular, Theorem 5 establishes nonasymptotic deterministic approximations for the bias $B_X(\hat{\beta}_\lambda; \beta)$ and variance $V_X(\hat{\beta}_\lambda; \beta)$. The error terms are uniform over the covariance matrix, and have nearly optimal dependence upon the sample size n . Indeed, a central-limit theorem heuristics suggests fluctuations of order $n^{-1/2}$.

As for the case of min-norm regression, Theorem 5 directly implies a characterization of the asymptotics of bias and variance of ridge regression. This statement is analogous to Theorem 6.

THEOREM 6. *Consider the setting of Theorem 5. Further assume $p/n \rightarrow \gamma \in (0, \infty)$, $\hat{H}_n \Rightarrow H$, $\hat{G}_n \Rightarrow G$. Define $\mathcal{B}_1(\lambda; H, G, \gamma)$ as in equation (43), with $\|\beta\|_2^2$ replaced by 1. Then, for any $\lambda > 0$, almost surely*

$$(48) \quad \frac{1}{\|\beta\|_2^2} B_X(\hat{\beta}_\lambda; \beta) \rightarrow \mathcal{B}_1(\lambda; H, G, \gamma), \quad V_X(\hat{\beta}_\lambda; \beta) \rightarrow \mathcal{V}(\lambda; H, \gamma).$$

The same conclusion holds if instead of Assumption 1(a), the coordinates of z , $(z_i)_{i \leq p}$ are i.i.d. and satisfy the conditions $\mathbb{E}z_i = 0$, $\mathbb{E}(z_i^2) = 1$, $\mathbb{E}(|z_i|^{4+\delta}) \leq C < \infty$.

6.1. Isotropic features. As a special case, we can consider the simple isotropic model that was already studied in Section 3. Very similar (though not identical) results can be found in Dicker [22], Dobriban and Wager [23].

COROLLARY 6. *Assume the conditions of Theorem 1 (well-specified model, isotropic features). Then for ridge regression estimator in (7) as $n, p \rightarrow \infty$, such that $p/n \rightarrow \gamma \in (0, \infty)$, it holds almost surely that*

$$(49) \quad R_X(\hat{\beta}_\lambda; \beta) \rightarrow r^2 \lambda^2 m'(-\lambda) + \sigma^2 \gamma (m(-\lambda) - \lambda m'(-\lambda)).$$

Here, $m(z)$ is given by equation (41), which in this case has the explicit solution $m(z) = [1 - \gamma - z - \sqrt{(1 - \gamma - z)^2 - 4\gamma z}]/(2\gamma z)$.

Furthermore, the limiting ridge risk is minimized at $\lambda^ = \sigma^2 \gamma / r^2$, in which case we have the simpler expression $R_X(\hat{\beta}_\lambda; \beta) \rightarrow \sigma^2 \gamma m(-\lambda^*)$.*

It is easy to recover the formulas in Theorem 1 as a limiting case of equation (49), by using the $z \rightarrow 0$ asymptotics $m(z) = (1 - \gamma)^{-1} + O(z)$ for $\gamma < 1$ and $m(z) = (1 - \gamma)^{-1} + O(z)$ for $\gamma < 1$ and $m(z) = -(\gamma - 1)/(\gamma z) + [(\gamma - 1)\gamma]^{-1} + O(z)$ for $\gamma > 1$.

Figures 7 and 8 compare the risk curves of min-norm least squares to those from optimally-tuned ridge regression, in the well-specified and misspecified settings, respectively. There are two important points to make. The first is that optimally-tuned ridge regression is seen to have strictly better asymptotic risk throughout, regardless of r^2 , γ , κ . This should not be a surprise, as by definition optimal tuning should yield better risk than min-norm least squares, which is the special case given by $\lambda \rightarrow 0^+$.

The second point is that, in this example, the limiting risk of optimally-tuned ridge regression appears to have a minimum around $\gamma = 1$, and this occurs closer and closer to $\gamma = 1$ as SNR grows. This behavior is interesting, especially because it is antipodal to that of the min-norm least squares risk, and leads us to very different suggestions for practical

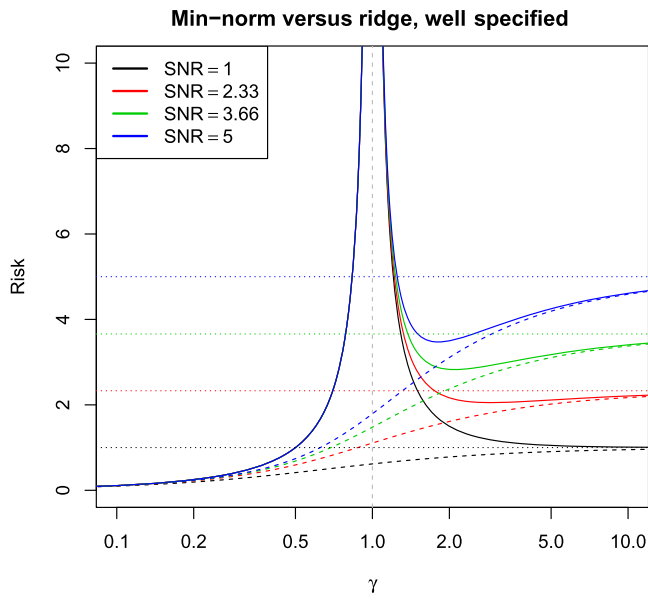


FIG. 7. Asymptotic risk curves for the min-norm least squares estimator in (10) as solid lines, and optimally-tuned ridge regression (from Theorem 5) as dashed lines. Here, r^2 varies from 1 to 5, and $\sigma^2 = 1$. The null risks are marked by the dotted lines.

usage for feature generators: in settings where we apply substantial ℓ_2 regularization (say, using CV tuning to mimic optimal tuning, which the next section shows to be asymptotically equivalent), it seems we want the complexity of the feature space to put us as close to the interpolation boundary ($\gamma = 1$) as possible.

As we will see, the behavior is rather different in the latent space model.

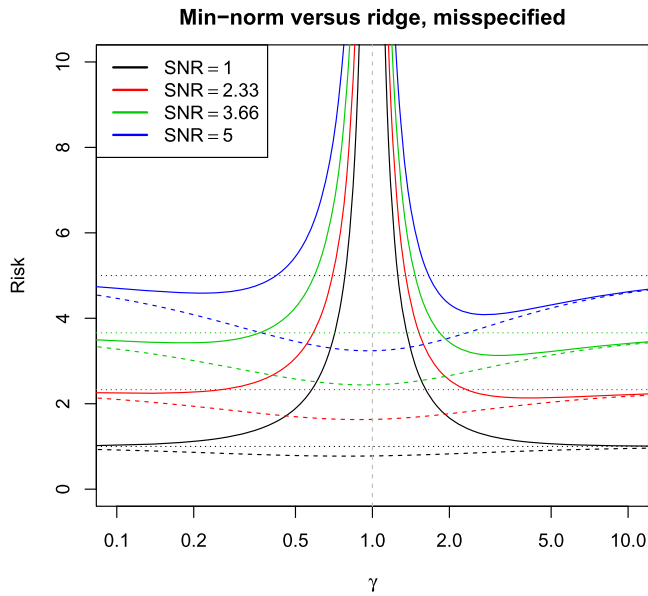


FIG. 8. Asymptotic risk curves for the min-norm least squares estimator in (28) as solid lines, and optimally-tuned ridge regression (from Theorem 5) as dashed lines, in the misspecified case, when the approximation bias has polynomial decay as in (27), with $a = 2$. Here, r^2 varies from 1 to 5, and $\sigma^2 = 1$. The null risks are marked by the dotted lines.

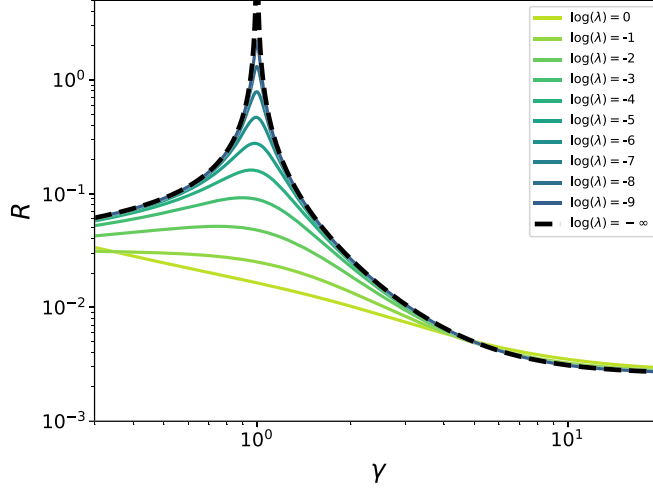


FIG. 9. Asymptotic risk as a function of the overparametrization ratio $\gamma = p/n$, for ridge regression in the latent space model of Section 6.2. Here, $n = 400$, $d = 20$, $\mu = 1$, $r_\theta = 1$, $\sigma_\xi = 0$, and each curve corresponds to a different value of the regularization λ . The dashed curve correspond to the min-norm interpolator (which coincides with the $\lambda \rightarrow 0$ limit of ridge regression).

6.2. Latent space model. As a special application, we consider the latent space model of Section 5.2. It is immediate to specialize equations (43) and (44) to this case. We omit giving explicit formulas for brevity, and instead plot the resulting curves for the prediction risk (test error).

In Figure 9, we plot the risk as a function of the overparametrization ratio $\gamma = p/n$ for several values of the regularization parameter λ (included the ridgeless limit $\lambda \rightarrow 0$). The setting here is analogous to the one of Figure 5. We observe several interesting phenomena:

1. Independently of λ in the probed range, the risk is minimized at large overparametrization $\gamma \gg 1$.
2. As expected, the divergence of the risk at the interpolation threshold $\gamma = 1$ is smoothed out by regularization, and the risk becomes a monotone decreasing function of γ when λ is large enough. Crucially, the optimal amount of regularization (corresponding to the lower envelope of these curves) results in a monotonically decreasing risk.
3. At large overparametrization, the optimal value of the regularization parameter is $\lambda \rightarrow 0$.

We confirm the last finding in Figure 10, which plots the risk as a function of λ : the optimal regularization is $\lambda \rightarrow 0$. Notice that this is the case despite the fact that the observations are noisy, namely $\sigma_\xi > 0$ strictly.

The optimality of $\lambda \rightarrow 0$ has a known intuitive explanation that is worth recalling here. Recall that ridge predictor at point x_0 takes the form

$$(50) \quad \hat{f}_\lambda(x_0) = \langle x_0, \hat{\beta}_\lambda \rangle = K(x_0, X)(K(X, X) + \lambda I)^{-1}y,$$

where we introduced the kernel matrix $K(X, X) = XX^T/n$, and the vector $K(x_0, X) = x_0^T X/n$. Consider the case in which the covariates X contain noise, as is our case, $X = \bar{X} + \eta U$ where $\bar{X} = ZW^T$ and $u_{ij} \sim N(0, 1)$. (While we are considering $\eta = 1$, it is instructive to regard the noise standard deviation as a parameter.) Then we might expect $K(X, X) \approx K(\bar{X}, \bar{X}) + \eta^2 U U^T \approx K(\bar{X}, \bar{X}) + \eta^2 I_n$. If this approximation holds, the noise acts as an extra ridge term, which can be sufficient to regularize the problem.

To the best of our knowledge, this argument was first presented by Webb [68] and, more explicitly, by Bishop [13]. Recently, Kobak et al. [42] elucidated its role in linear regression,

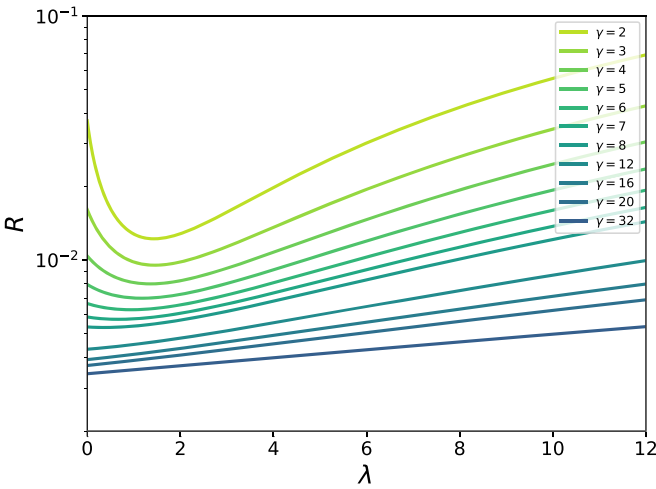


FIG. 10. Asymptotic risk as a function of the regularization parameter λ , for ridge regression in the latent space model of Section 6.2. Here, $n = 400$, $d = 20$, $\mu = 1$, $r_\theta = 1$, $\sigma_\xi = 0.1$, and each curve corresponds to a different value of the overparametrization ratio $\gamma = p/n$.

establishing several of its consequences. In a parallel line of work, a closely related idea has recently emerged in the analysis of kernel methods in high dimension [26, 47].

7. Cross-validation. We analyze the effect of using cross-validation to choose the tuning parameter in ridge regression. In short, we find that choosing the ridge tuning parameter to minimize the leave-one-out cross-validation error leads to the same asymptotic risk as the optimally-tuned ridge estimator. The next subsection gives the details; the following subsection presents a new “shortcut formula” for leave-one-out cross-validation in the overparametrized regime, for min-norm least squares, akin to the well-known formula for underparametrized least squares and ridge regression. We refer to [20, 32] for background on CV and GCV, and to [5] for a more recent review.

7.1. Limiting behavior of CV tuning. Given the ridge regression solution $\hat{\beta}_\lambda$ in (7), trained on (x_i, y_i) , $i = 1, \dots, n$, denote by $\hat{f}_\lambda$ the corresponding ridge predictor, defined as $\hat{f}_\lambda(x) = x^T \hat{\beta}_\lambda$ for $x \in \mathbb{R}^p$. Additionally, for each $i = 1, \dots, n$, denote by $\hat{f}_\lambda^{-i}$ the ridge predictor trained on all but i th data point (x_i, y_i) .⁵ Recall that the *leave-one-out cross-validation* (leave-one-out CV, or simply CV) error of the ridge solution at a tuning parameter value λ is

(51)
$$\text{CV}_n(\lambda) = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{f}_\lambda^{-i}(x_i))^2.$$

We typically view this as an estimate of the out-of-sample prediction error $\mathbb{E}(y_0 - x_0^T \hat{\beta}_\lambda)^2$, where the expectation is taken over everything that is random: the training data (x_i, y_i) , $i = 1, \dots, n$ used to fit $\hat{\beta}_\lambda$, as well as the independent test point (x_0, y_0) . Note also that, when we observe training data from the model (1), (2), and when (x_0, y_0) is drawn independently according to the same process, we have the relationship

$$\mathbb{E}(y_0 - x_0^T \hat{\beta}_\lambda)^2 = \sigma^2 + \mathbb{E}(x_0^T \beta - x_0^T \hat{\beta}_\lambda)^2 = \sigma^2 + \mathbb{E}[R_X(\hat{\beta}_\lambda; \beta)],$$

⁵To be precise, this is $\hat{f}^{-i}(x) = x^T (X_{-i}^T X_{-i} + n\lambda I)^{-1} X_{-i}^T y_{-i}$, where X_{-i} denotes X with the i th row removed, and y_{-i} denotes y with the i th component removed. Arguably, it may seem more natural to replace the factor of n here by a factor of $n - 1$; we leave the factor of n as is because it simplifies the presentation in what follows, but we remark that the same asymptotic results would hold with $n - 1$ in place of n .

where $R_X(\hat{\beta}_\lambda; \beta) = \mathbb{E}[(x_0^T \beta - x_0^T \hat{\beta}_\lambda)^2 | X]$ is the conditional prediction risk, which has been our focus throughout.

Recomputing the leave-one-out predictors $\hat{f}_\lambda^{-i}$, $i = 1, \dots, n$ can be burdensome, especially for large n . Importantly, there is a well-known “shortcut formula” that allows us to express the leave-one-out CV error (51) as a weighted average of the training errors,

$$(52) \quad \text{CV}_n(\lambda) = \frac{1}{n} \sum_{i=1}^n \left(\frac{y_i - \hat{f}_\lambda(x_i)}{1 - (S_\lambda)_{ii}} \right)^2,$$

where $S_\lambda = X(X^T X + n\lambda)^{-1} X^T$ is the ridge smoother matrix. This identity is an immediate consequence of the Sherman–Morrison–Woodbury formula. In the next subsection, we give an extension to the case $\lambda = 0$ and $\text{rank}(X) = n$, that is, to min-norm least squares.

The next result shows that, for isotropic features, the CV error of a ridge estimator converges almost surely to its prediction error. The focus on isotropic features is only for simplicity: a more general analysis is possible but is not pursued here. The proof, given in Appendix A.4.6, relies on the shortcut formula (52). In the proof, we actually first analyze generalized cross-validation (GCV), which turns out to be somewhat of an easier calculation (see the proof for details on the precise form of GCV), and then relate leave-one-out CV to GCV.

THEOREM 7. *Assume the a isotropic prior, namely $\mathbb{E}(\beta) = 0$, $\text{Cov}(\beta) = r^2 I_p/p$, and the data model (1), (2). Assume that $x \sim P_x$ has i.i.d. entries with zero mean, unit variance, and a finite moment of order $4 + \eta$, for some $\eta > 0$. Then for the CV error (51) of the ridge estimator in (7) with tuning parameter $\lambda > 0$, as $n, p \rightarrow \infty$, with $p/n \rightarrow \gamma \in (0, \infty)$, it holds almost surely that*

$$\text{CV}_n(\lambda) - \sigma^2 \rightarrow \sigma^2 \gamma (m(-\lambda) - \lambda(1 - \alpha\lambda)m'(-\lambda)),$$

where $m(z)$ denotes the Stieltjes transform of the Marchenko–Pastur law F_γ (as in Corollary 6), and $\alpha = r^2/(\sigma^2 \gamma)$. Observe that the right-hand side is the asymptotic risk of ridge regression from Theorem 5. Moreover, the above convergence is uniform over compact intervals excluding zero. Thus if λ_1, λ_2 are constants with $0 < \lambda_1 \leq \lambda^* \leq \lambda_2 < \infty$, where $\lambda^* = 1/\alpha$ is the asymptotically optimal ridge tuning parameter value, and we define $\lambda_n = \arg \min_{\lambda \in [\lambda_1, \lambda_2]} \text{CV}_n(\lambda)$, then the expected risk of the CV-tuned ridge estimator $R_X(\hat{\beta}) := \mathbb{E}_\beta R_X(\hat{\beta}; \beta)$ satisfies, almost surely

$$R_X(\hat{\beta}_{\lambda_n}) \rightarrow \sigma^2 \gamma m(-1/\alpha),$$

with the right-hand side above being the asymptotic risk of optimally-tuned ridge regression. Further, the exact same set of results holds for GCV.

Similar results were obtained for various linear smoothers in Li [45, 46], for the lasso in the high-dimensional (proportional) asymptotics in Miolane and Montanari [51], and for general smooth penalized estimators in Xu et al. [70]. The latter paper covers ridge regression as a special case, and gives more precise results (convergence rates), but assumes more restrictive conditions. After submission of this paper, consistency of CV and GVC was proved in a significantly more general setting in [54], in particular dispensing with the assumption of random β .

The key implication of Theorem 7, in the context of the current paper and its central focus, is that the CV-tuned or GCV-tuned ridge estimator has the same asymptotic performance as the optimally-tuned ridge estimator. In other words, the ridge curves in Figures 1, 7 and 8 can be alternatively viewed as the asymptotic risk of ridge under CV tuning.

7.2. Shortcut formula for ridgeless CV. We extend the leave-one-out CV shortcut formula (52) to work when $p > n$ and $\lambda = 0+$, that is, for min-norm least squares. In this case, both the numerator and denominator are zero in each summand of (52). To circumvent this, we can use the so-called “kernel trick” to rewrite the ridge regression solution (7) with $\lambda > 0$ as

$$(53) \quad \hat{\beta}_\lambda = X^T (XX^T + n\lambda I)^{-1} y.$$

Using this expression, the shortcut formula for leave-one-out CV in (52) can be rewritten as

$$\text{CV}_n(\lambda) = \frac{1}{n} \sum_{i=1}^n \frac{[(XX^T + n\lambda I)^{-1} y]_i^2}{[(XX^T + n\lambda I)^{-1}]_{ii}^2}.$$

Taking $\lambda \rightarrow 0^+$ yields the shortcut formula for leave-one-out CV in min-norm least squares (assuming without a loss of generality that $\text{rank}(X) = n$),

$$(54) \quad \text{CV}_n(0) = \frac{1}{n} \sum_{i=1}^n \frac{[(XX^T)^{-1} y]_i^2}{[(XX^T)^{-1}]_{ii}^2}.$$

In fact, the exact same arguments given here still apply when we replace XX^T by a positive definite kernel matrix K (i.e., $K_{ij} = k(x_i, x_j)$ for each $i, j = 1, \dots, n$, where k is a positive definite kernel function), in which case (54) gives a shortcut formula for leave-one-out CV in kernel ridgeless regression (the limit in kernel ridge regression as $\lambda \rightarrow 0^+$). We also remark that, when we include an unpenalized intercept in the model, in either the linear or kernelized setting, the shortcut formula (54) still applies with XX^T or K replaced by their doubly-centered (row- and column-centered) versions, and the matrix inverses replaced by pseudoinverses.

8. Nonlinear model. Our analysis in the previous sections assumed $x_i = \Sigma^{1/2} z_i$, with z_i a vector with independent entries.

In this section, we test universality on one simple example. We observe data as in (1), (2), but now $x_i = \varphi(Wz_i) \in \mathbb{R}^p$, where $z_i \in \mathbb{R}^d$ has i.i.d. entries from $N(0, 1)$, for $i = 1, \dots, n$. Also, $W \in \mathbb{R}^{p \times d}$ has i.i.d. entries from $N(0, 1/d)$. Finally, $\varphi : \mathbb{R} \rightarrow \mathbb{R}$ is an activation function acting entrywise on vectors.

We first consider the case of purely nonlinear activations, namely activation functions that are uncorrelated with linear functions: $\mathbb{E}\{\varphi(G)\} = \mathbb{E}\{G\varphi(G)\} = 0$ for $G \sim N(0, 1)$. In this case, the second-order statistics of the features x_i match the ones of the isotropic model and the same happen for the asymptotic of the risk. We then consider more general activations, and show that the asymptotic variance depends on the activation through the value of $c_1 = \mathbb{E}\{G\varphi(G)\}^2$.

8.1. Limiting risk for purely nonlinear activations. Notice that, conditionally on W , the vectors $x_i = \varphi(Wz_i)$, $i \leq n$ are independent. However, they do not have independent coordinates. For instance, if $\varphi(t) = at^2 + b$, we can reconstruct z_i from the first $2d$ coordinates of x_i and, therefore, the remaining $p - 2d$ coordinates of x_i are a function of the first $2d$.

Nevertheless, the next theorem shows that if φ is purely nonlinear (in the sense that $\mathbb{E}\{\varphi(G)\} = \mathbb{E}\{G\varphi(G)\} = 0$), then the feature matrix X behaves “as if” it had i.i.d. entries, in that the asymptotic bias and variance are exactly as in the linear isotropic case; recall equation (10). In other words, this theorem provides a rigorous confirmation of the universality hypothesis stated in the [Introduction](#).

THEOREM 8. Assume the model (1), (2), where each $x_i = \varphi(Wz_i) \in \mathbb{R}^p$, for $z_i \in \mathbb{R}^d$ having i.i.d. entries from $N(0, 1)$, $W \in \mathbb{R}^{p \times d}$ having i.i.d. entries from $N(0, 1/d)$ (with W independent of z_i), and is φ an activation function that acts componentwise. Assume that $|\varphi(x)| \leq c_0(1 + |x|)^{c_0}$ for a constant $c_0 > 0$. Also, for $G \sim N(0, 1)$, assume that the following standardization conditions hold: $\mathbb{E}[\varphi(G)] = 0$ and $\mathbb{E}[\varphi(G)^2] = 1$, $\mathbb{E}[G\varphi(G)] = 0$. Consider the limit $n, p, d \rightarrow \infty$, with $p/n \rightarrow \gamma$ and $d/p \rightarrow \psi \in (0, \infty)$.

Then for $\gamma > 1$, the variance satisfies, almost surely

$$\lim_{\lambda \rightarrow 0^+} \lim_{n, p, d \rightarrow \infty} V_X(\hat{\beta}_\lambda; \beta) = \frac{\sigma^2}{\gamma - 1},$$

which is precisely as in the case of linear isotropic features; recall Theorem 1. Also, under a isotropic prior, namely $\mathbb{E}(\beta) = 0$, $\text{Cov}(\beta) = r^2 I_p/p$, the Bayes bias $B_X(\hat{\beta}_\lambda) := \mathbb{E}_\beta B_X(\hat{\beta}_\lambda; \beta)$ satisfies, almost surely

$$\lim_{\lambda \rightarrow 0^+} \lim_{n, p, d \rightarrow \infty} B_X(\hat{\beta}_\lambda) = \begin{cases} 0 & \text{for } \gamma < 1, \\ r^2(1 - 1/\gamma) & \text{for } \gamma > 1, \end{cases}$$

which is again as in the case of linear isotropic features; recall Theorem 1.

The proof of Theorem 8 is lengthy and will be sketched shortly. We notice that the definition of $V_X(\hat{\beta}_\lambda; \beta)$ and $B_X(\hat{\beta}_\lambda)$ is conditional on X . In fact, we will prove that the stated limits hold asymptotically almost surely, conditionally both on W and on the covariates x_i .

The origin of the conditions $\mathbb{E}\{\varphi(G)\} = \mathbb{E}\{G\varphi(G)\} = 0$ can be easily explained (throughout $G \sim N(0, 1)$). In summary, these conditions ensure that the first- and second-order statistics of $x_i = \varphi(Wz_i)$ approximately match those of the isotropic model. To illustrate this point, let $i \neq j$, and assume that the corresponding rows of W (denoted by w_i^T and w_j^T) have unit norm (this will only be approximately true, but simplifies our explanation). We then have $\mathbb{E}_z\{\varphi(w_i^T z_1)\} = \mathbb{E}_z\{\varphi(w_j^T z_1)\} = \mathbb{E}\{\varphi(G)\} = 0$. Further,

$$(55) \quad \mathbb{E}_z\{x_{1,i}x_{1,j}|W\} = \mathbb{E}_z\{\varphi(w_i^T z_1)\varphi(w_j^T z_1)|W\} = \mathbb{E}\{\varphi(G_1)\varphi(G_2)\},$$

where G_1, G_2 are jointly Gaussian with unit variance and covariance $w_i^T w_j$. Denoting by $\varphi(x) = \sum_{k \geq 0} \lambda_k(\varphi) h_k(x)$ the decomposition of φ into orthonormal Hermite polynomials, we thus obtain

$$(56) \quad \mathbb{E}\{x_{1,i}x_{1,j}|W\} = \sum_{k=0}^{\infty} \lambda_k^2(\varphi) (w_i^T w_j)^k,$$

Since $\lambda_0(\varphi) = \mathbb{E}\{\varphi(G)\} = 0$, $\lambda_1(\varphi) = \mathbb{E}\{G\varphi(G)\} = 0$, we have $\mathbb{E}\{x_{1,i}x_{1,j}|W\} = O((w_i^T w_j)^2) = O(1/d)$. In other words, the population covariance $\mathbb{E}\{x_1 x_1^T | W\}$ has small entries out-of diagonal, and in fact $\|\mathbb{E}\{x_1 x_1^T | W\} - I_p\|_{\text{op}} = o_P(1)$ [26].

Even if the nonlinear model $x_i = \varphi(Wz_i)$ matches the second-order population statistics of the isotropic model, it is far from obvious that the asymptotics of the risk is the same. Indeed the coordinates of vector x_i are highly dependent. Theorem 8 confirms that—despite dependence—the risk is asymptotically the same, thus providing a concrete example of the general universality phenomenon.

Figure 11 compares the asymptotic risk curve from Theorem 8 to that computed by simulation, using an activation function $\varphi_{\text{abs}}(t) = a(|t| - b)$, where $a = \sqrt{\pi/(\pi - 2)}$ and $b = \sqrt{2/\pi}$ are chosen to meet the standardization conditions. This activation function is purely nonlinear, that is, it satisfies $\mathbb{E}[G\varphi_{\text{abs}}(G)] = 0$ for $G \sim N(0, 1)$, by symmetry. Again, the agreement between finite-sample and asymptotic risks is excellent. Notice in particular that, as predicted by Theorem 8, the risk depends only on p/n and not on d/n .

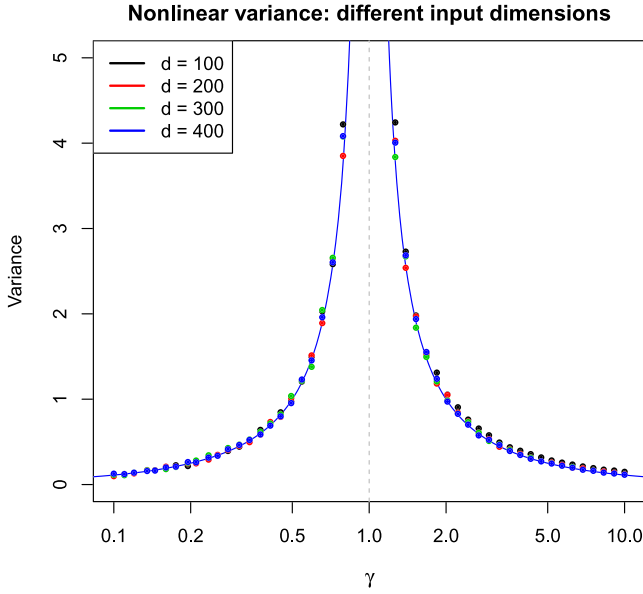


FIG. 11. Asymptotic variance curves for the min-norm least squares estimator in the nonlinear feature model (from Theorem 8), for the purely nonlinear activation φ_{abs} . Here $\sigma^2 = 1$, and the points are finite-sample risks, with $n = 200$, $p = \lceil \gamma n \rceil$, over various values of γ , and varying input dimensions: $d = 100$ in black, $d = 200$ in red, $d = 300$ in green, and $d = 400$ in black. As before, the features used for finite-sample calculations are $X = \varphi(ZW^T)$, where Z has i.i.d. $N(0, 1)$ entries and W has i.i.d. $N(0, 1/d)$ entries.

8.2. *Limiting variance for general activations.* Consider now the case of a general activation with vanishing mean $\mathbb{E}\{\varphi(G)\} = 0$, but nonvanishing linear component $\mathbb{E}\{G\varphi(G)\}^2 = c_1$, and normalized so that $\mathbb{E}\{\varphi(G)^2\} = 1$. Following the same argument as in the last section, we obtain the following approximation for the conditional covariance of x_i given W :

$$(57) \quad \|\mathbb{E}\{x_1 x_1^T | W\} - \tilde{\Sigma}_{X|W}\|_{\text{op}} = o_P(1), \quad \tilde{\Sigma}_{X|W} = (1 - c_1)I_p + c_1 W W^T.$$

In other words, the conditional covariance is well approximated by the covariance of the the latent space model of Sections 5.2 and 6.2. Let us emphasize that the conditional covariance is what matters (not the unconditional one) because the x_i 's are independent only conditional on W .

In Appendix A.5, we derive the asymptotic variance for this model. The next result confirms once more the universality scenario outlined above. In words, the asymptotic variance of the nonlinear model $x_i = \varphi(Wz_i)$ coincides with the asymptotic variance of the corresponding linear model $x_i \sim N(0, \Sigma_W)$.

THEOREM 9. *In the setting of Theorem 8, assume $\mathbb{E}[\varphi(G)] = 0$ and $\mathbb{E}[\varphi(G)^2] = 1$, $\mathbb{E}[G\varphi(G)]^2 = c_1$, and let $V_X(\hat{\beta}_\lambda; \beta)$ denote the corresponding variance of ridge regression.*

Further, consider a different problem with $\tilde{x}_i \sim N(0, \tilde{\Sigma}_{X|W})$, $\tilde{\Sigma}_{X|W} := (1 - c_1)I_p + c_1 W W^T$, and denote by $V_{\tilde{X}}(\hat{\beta}_\lambda; \beta)$ denote the corresponding variance of ridge regression. Assume $p, n, d, \rightarrow \infty$ with $p/n \rightarrow \gamma \in (1, \infty$, and $d/p \rightarrow \psi \in (0, \infty)$. Then we have, almost surely

$$\lim_{\lambda \rightarrow 0^+} \lim_{n, p, d \rightarrow \infty} V_X(\hat{\beta}_\lambda; \beta) = \lim_{\lambda \rightarrow 0^+} \lim_{n, p, d \rightarrow \infty} V_{\tilde{X}}(\hat{\beta}_\lambda; \beta).$$

REMARK 5. After the last result was proved and a preprint posted online, the techniques developed here were significantly sharpened and generalized in [49], which proved in particular that the same universality result holds for the bias term as well. We notice that both

Theorem 8 and Theorems 9 as well as its generalization in [49] assume a particularly simple model for the latent covariates z_i . While this simple model simplifies the proof, the result is likely to generalize to other models, for example, z_i with independent sub-Gaussian entries.

REMARK 6. Notice that the assumption of isotropic W in Theorem 8 and Theorem 9 is realistic as this is the standard choice of random features models. On the other hand, it is an interesting research question to which extent these results generalize to other distributions of z_i , for example, $z_i \sim N(0, \Sigma_Z)$. We believe that the results obtained here extend to that case as well as long as $p/n \rightarrow \gamma \in (1, \infty)$, and $d/p \rightarrow \psi \in (0, \infty)$ and Σ_Z has a positive fraction of eigenvalues of the same order as its largest eigenvalue (as requested for Σ in Assumption 1). In that case of course, one has to replace the above formula for $\tilde{\Sigma}_{X|W}$ by $\tilde{\Sigma}_{X|W} = \mathbb{E}_z \{\varphi(Wz)\varphi(Wz)^T\}$.

8.3. *Proof outline for Theorem 8.* We define $\gamma_n = p/n$ and $\psi_n = d/p$. Recall that as $n, p, d \rightarrow \infty$, we have $\gamma_n \rightarrow \gamma$ and $\psi_n \rightarrow \psi$. To reduce notational overhead, we will generally drop the subscripts from γ_n, ψ_n , writing these simply as γ, ψ , since their meanings should be clear from the context. Let $N = p + n$ and define the symmetric matrix $A(s) \in \mathbb{R}^{N \times N}$, for $s \geq 0$, with the block structure:

$$(58) \quad A(s) = \begin{bmatrix} sI_p & \frac{1}{\sqrt{n}}X^T \\ \frac{1}{\sqrt{n}}X & 0_n \end{bmatrix},$$

where $I_p \in \mathbb{R}^{p \times p}$ and $0_n \in \mathbb{R}^{n \times n}$ are the identity and zero matrix, respectively. As we will see, this matrix allows to construct the traces of interest by taking suitable derivatives of its resolvent.

We introduce the following resolvents (as usual, these are defined for $\text{Im}(\xi) > 0$ and by analytic continuation, whenever possible, for $\text{Im}(\xi) = 0$):

$$\begin{aligned} m_{1,n}(\xi, s) &= \mathbb{E}\{(A(s) - \xi I_N)^{-1}_{1,1}\} = \mathbb{E}M_{1,n}(\xi, s), \\ M_{1,n}(\xi, s) &= \frac{1}{p} \text{Tr}_{[1,p]} \{(A(s) - \xi I_N)^{-1}\}, \\ m_{2,n}(\xi, s) &= \mathbb{E}\{(A(s) - \xi I_N)^{-1}_{p+1,p+1}\} = \mathbb{E}M_{2,n}(\xi, s), \\ M_{2,n}(\xi, s) &= \frac{1}{n} \text{Tr}_{[p+1,p+n]} \{(A(s) - \xi I_N)^{-1}\}. \end{aligned}$$

Here and henceforth, we write $[i, j] = \{i + 1, \dots, i + j\}$ for integers i, j . We also write $M_{ij}^{-1} = (M^{-1})_{ij}$ for a matrix M , and $\text{Tr}_S(M) = \sum_{i \in S} M_{ii}$ for a subset S . The equalities in the first and third lines above follow by invariance of the distribution of $A(s)$ under permutations of $[1, p]$ and $[p + 1, p + n]$. Whenever clear from the context, we will omit the arguments from block matrix and resolvents, and write $A = A(s)$, $m_{1,n} = m_{1,n}(\xi, s)$ and $m_{2,n} = m_{2,n}(\xi, s)$.

The next lemma characterizes the asymptotics of $m_{1,n}, m_{2,n}$.

LEMMA 3. Assume the conditions of Theorem 8. Consider $\text{Im}(\xi) > 0$ or $\text{Im}(\xi) = 0$, $\text{Re}(\xi) < 0$, with $s \geq t \geq 0$. Let m_1 and m_2 be the unique solutions of the following quadratic equations:

$$(59) \quad m_2 = (-\xi - \gamma m_1)^{-1}, \quad m_1 = (-\xi - s - m_2)^{-1},$$

subject to the condition of being analytic functions for $\text{Im}(z) > 0$, and satisfying $|m_1(z, s)|, |m_2(z, s)| \leq 1/\text{Im}(z)$ for $\text{Im}(z) > C$ (with C a sufficiently large constant). Then, as $n, p, d \rightarrow \infty$, such that $p/n \rightarrow \gamma$ and $d/p \rightarrow \psi$, we have almost surely (and in L^1),

$$(60) \quad \lim_{n,p,d \rightarrow \infty} M_{1,n}(\xi, s) = m_1(\xi, s),$$

$$(61) \quad \lim_{n,p,d \rightarrow \infty} M_{2,n}(\xi, s) = m_2(\xi, s).$$

The proof of this lemma is given in Appendix A.5.2. As a corollary of the above, we obtain that the asymptotical empirical spectral distribution of the empirical covariance $\hat{\Sigma} = X^T X/n$ matches the one for the independent entries model, and is hence given by the Marchenko–Pastur law (a result already obtained in Pennington and Worah [55]). We state this formally using the Stieltjes transform

$$(62) \quad R_n(z) = \frac{1}{p} \text{Tr}((\hat{\Sigma} - zI_p)^{-1}).$$

COROLLARY 7. Assume the conditions of Theorem 8. Consider $\text{Im}(z) > 0$. As $n, p, d \rightarrow \infty$, with $p/n \rightarrow \gamma$ and $d/p \rightarrow \psi$, we have (almost surely and in L^1) $R_n(\xi) \rightarrow r(\xi)$ where r is nonrandom and coincides with the Stieltjes transform of the Marchenko–Pastur law, namely

$$(63) \quad r(z) = \frac{1 - \gamma - z - \sqrt{(1 - \gamma - z)^2 - 4\gamma z}}{2\gamma z}.$$

We refer to Appendix A.5.4 for a proof of this corollary. The next lemma connects the above resolvents computed in Lemma 3 to the variance of min-norm least squares, hence completing our proof outline.

LEMMA 4. Assume the conditions of Theorem 8. Let m_1, m_2 be the asymptotic resolvents given in Lemma 3. Define

$$m(\xi, s) = \gamma m_1(\xi, s) + m_2(\xi, s).$$

Then for $\gamma \neq 1$, $\partial_x m(\xi, x)|_{x=0}$ as a simple pole at $\xi = 0$, and hence admits a Taylor–Laurent expansion around $\xi = 0$, whose coefficients will be denoted by D_{-1}, D_0 ,

$$(64) \quad -\partial_x m(\xi, x)|_{x=0} = \frac{D_{-1}}{\xi^2} + D_0 + O(\xi^2).$$

Here, each coefficient is a function of γ, ψ : $D_{-1} = D_{-1}(\gamma, \psi)$, $D_0 = D_0(\gamma, \psi)$. Furthermore, for the ridge regression estimator $\hat{\beta}_\lambda$ in (7), as $n, p, d \rightarrow \infty$, such that $p/n \rightarrow \gamma \in (0, \infty)$, $d/p \rightarrow \psi \in (0, 1)$, the following ridgeless limit holds almost surely:

$$\lim_{\lambda \rightarrow 0^+} \lim_{n,p,d \rightarrow \infty} V_X(\hat{\beta}_\lambda; \beta) = D_0.$$

The proof of this lemma can be found in Appendix A.5.3. Theorem 8 follows by evaluating the formula in Lemma 4, by using the result of Lemma 3. We refer to the Appendix in the Supplementary Material for details [36].

Acknowledgments. The authors are grateful to Brad Efron, Rob Tibshirani and Larry Wasserman for inspiring us to work on this in the first place. RT sincerely thanks Edgar Dobriban for many helpful conversations about the random matrix theory literature, in particular, the literature trail leading up to Proposition 2, and the reference to Rubio and Mestre [60]. We are also grateful to Dmitry Kobak for bringing to our attention [42] and clarifying the implications of this work on our setting.

Funding. TH was partially supported by NSF Grants DMS-1407548, NSF IIS-1837931 and NIH 5R01-EB001988-21.

AM was partially supported by NSF Grants DMS-1613091, NSF CCF-1714305, NSF IIS-1741162 and ONR N00014-18-1-2729.

RT was partially supported by NSF Grant DMS-1554123.

SUPPLEMENTARY MATERIAL

Supplement to “Surprises in high-dimensional ridgeless least squares interpolation” (DOI: [10.1214/21-AOS2133SUPP](https://doi.org/10.1214/21-AOS2133SUPP); .pdf). Supplementary information.

REFERENCES

- [1] ADLAM, B. and PENNINGTON, J. (2020). The neural tangent kernel in high dimensions: Triple descent and a multi-scale theory of generalization. Available at <http://proceedings.mlr.press/v119/adlam20a/adlam20a.pdf>.
- [2] ADVANI, M. S. and SAXE, A. M. (2017). High-dimensional dynamics of generalization error in neural networks. Available at [arXiv:1710.03667](https://arxiv.org/abs/1710.03667).
- [3] ALI, A., KOLTER, J. Z. and TIBSHIRANI, R. J. (2019). A continuous-time view of early stopping for least squares. *Int. Conf. Artif. Intell. Stat.* **22**.
- [4] ALLEN-ZHU, Z., LI, Y. and SONG, Z. (2019). A convergence theory for deep learning via over-parameterization. International Conference on Machine Learning. PMLR. Available at <http://proceedings.mlr.press/v97/allen-zhu19a/allen-zhu19a.pdf>.
- [5] ARLOT, S. and CELISSE, A. (2010). A survey of cross-validation procedures for model selection. *Stat. Surv.* **4** 40–79. [MR2602303 https://doi.org/10.1214/09-SS054](https://doi.org/10.1214/09-SS054)
- [6] BARTLETT, P. L., LONG, P. M., LUGOSI, G. and TSIGLER, A. (2020). Benign overfitting in linear regression. *Proc. Natl. Acad. Sci. USA* **117** 30063–30070. [MR4263288 https://doi.org/10.1073/pnas.1907378117](https://doi.org/10.1073/pnas.1907378117)
- [7] BARTLETT, P. L., MONTANARI, A. and RAKHLIN, A. (2021). Deep learning: A statistical viewpoint. *Acta Numer.* **30** 87–201. [MR4295218 https://doi.org/10.1017/S0962492921000027](https://doi.org/10.1017/S0962492921000027)
- [8] BELKIN, M., HSU, D., MA, S. and MANDAL, S. (2019). Reconciling modern machine-learning practice and the classical bias-variance trade-off. *Proc. Natl. Acad. Sci. USA* **116** 15849–15854. [MR3997901 https://doi.org/10.1073/pnas.1903070116](https://doi.org/10.1073/pnas.1903070116)
- [9] BELKIN, M., HSU, D. and XU, J. (2020). Two models of double descent for weak features. *SIAM J. Math. Data Sci.* **2** 1167–1180. [MR4186534 https://doi.org/10.1137/20M11336072](https://doi.org/10.1137/20M11336072)
- [10] BELKIN, M., MA, S. and MANDAL, S. (2018). To understand deep learning we need to understand kernel learning. 22nd International Conference on Artificial Intelligence and Statistics. PMLR. Available at <http://proceedings.mlr.press/v80/belkin18a/belkin18a.pdf>.
- [11] BELKIN, M., RAKHLIN, A. and TSYBAKOV, A. B. (2018). Does data interpolation contradict statistical optimality? Available at [arXiv:1806.09471](https://arxiv.org/abs/1806.09471).
- [12] BING, X., BUNEA, F., STRIMAS-MACKEY, S. and WEGKAMP, M. (2021). Prediction under latent factor regression: Adaptive PCR, interpolating predictors and beyond. *J. Mach. Learn. Res.* **22** 177. [MR4318533 https://doi.org/10.22405/2226-8383-2021-22-1-177-187](https://doi.org/10.22405/2226-8383-2021-22-1-177-187)
- [13] BISHOP, C. M. (1995). Training with noise is equivalent to Tikhonov regularization. *Neural Comput.* **7** 108–116.
- [14] BUNEA, F., STRIMAS-MACKEY, S. and WEGKAMP, M. (2020). Interpolating predictors in high-dimensional factor regression. Available at [arXiv:2002.02525](https://arxiv.org/abs/2002.02525).
- [15] CHATTERJI, N. S. and LONG, P. M. (2021). Finite-sample analysis of interpolating linear classifiers in the overparameterized regime. *J. Mach. Learn. Res.* **22** 129. [MR4279780](https://doi.org/10.1137/S1064827596304010)
- [16] CHEN, S. S., DONOHO, D. L. and SAUNDERS, M. A. (1998). Atomic decomposition by basis pursuit. *SIAM J. Sci. Comput.* **20** 33–61. [MR1639094 https://doi.org/10.1137/S1064827596304010](https://doi.org/10.1137/S1064827596304010)
- [17] CHENG, X. and SINGER, A. (2013). The spectrum of random inner-product kernel matrices. *Random Matrices Theory Appl.* **2** 1350010. [MR3149440 https://doi.org/10.1142/S201032631350010X](https://doi.org/10.1142/S201032631350010X)
- [18] CHIZAT, L. and BACH, F. (2018). On the global convergence of gradient descent for over-parameterized models using optimal transport. *Adv. Neural Inf. Process. Syst.* **31**.
- [19] CHIZAT, L. and BACH, F. (2019). A note on lazy training in supervised differentiable programming. *Advances in Neural Information Processing Systems* **32** 2933–2943.

- [20] CRAVEN, P. and WAHBA, G. (1978/79). Smoothing noisy data with spline functions. Estimating the correct degree of smoothing by the method of generalized cross-validation. *Numer. Math.* **31** 377–403. [MR0516581 https://doi.org/10.1007/BF01404567](https://doi.org/10.1007/BF01404567)
- [21] DAUBECHIES, I. (1988). Time-frequency localization operators: A geometric phase space approach. *IEEE Trans. Inf. Theory* **34** 605–612. [MR0966733 https://doi.org/10.1109/18.9761](https://doi.org/10.1109/18.9761)
- [22] DICKER, L. H. (2016). Ridge regression and asymptotic minimax estimation over spheres of growing dimension. *Bernoulli* **22** 1–37. [MR3449775 https://doi.org/10.3150/14-BEJ609](https://doi.org/10.3150/14-BEJ609)
- [23] DOBRIBAN, E. and WAGER, S. (2018). High-dimensional asymptotics of prediction: Ridge regression and classification. *Ann. Statist.* **46** 247–279. [MR3766952 https://doi.org/10.1214/17-AOS1549](https://doi.org/10.1214/17-AOS1549)
- [24] DU, S. S., LEE, J. D., LI, H., WANG, L. and ZHAI, X. (2018). Gradient descent finds global minima of deep neural networks. International conference on machine learning. PMLR.
- [25] DU, S. S., ZHAI, X., POCZOS, B. and SINGH, A. (2018). Gradient descent provably optimizes over-parameterized neural networks. Available at [arXiv:1810.02054](https://arxiv.org/abs/1810.02054).
- [26] EL KAROUI, N. (2010). The spectrum of kernel random matrices. *Ann. Statist.* **38** 1–50. [MR2589315 https://doi.org/10.1214/08-AOS648](https://doi.org/10.1214/08-AOS648)
- [27] FAN, Z. and MONTANARI, A. (2019). The spectral norm of random inner-product kernel matrices. *Probab. Theory Related Fields* **173** 27–85. [MR3916104 https://doi.org/10.1007/s00440-018-0830-4](https://doi.org/10.1007/s00440-018-0830-4)
- [28] GEIGER, M., JACOT, A., SPIGLER, S., GABRIEL, F., SAGUN, L., D’ASCOLI, S., BIROLI, G., HONGLER, C. and WYART, M. (2020). Scaling description of generalization with number of parameters in deep learning. *J. Stat. Mech. Theory Exp.* **2023401**. [MR4137317 https://doi.org/10.1088/1742-5468/ab633c](https://doi.org/10.1088/1742-5468/ab633c)
- [29] GERACE, F., LOUREIRO, B., KRZAKALA, F., MÉZARD, M. and ZDEBOROVÁ, L. (2020). Generalisation error in learning with random features and the hidden manifold model. International Conference on Machine Learning. PMLR.
- [30] GHORBANI, B., MEI, S., MISIAKIEWICZ, T. and MONTANARI, A. (2021). Linearized two-layers neural networks in high dimension. *Ann. Statist.* **49** 1029–1054. [MR4255118 https://doi.org/10.1214/20-aos1990](https://doi.org/10.1214/20-aos1990)
- [31] GOLDT, S., REEVES, G., MEZARD, M., KRZAKALA, F. and ZDEBOROVÁ, L. (2020). The gaussian equivalence of generative models for learning with two-layer neural networks. Available at [arXiv:2006.14709](https://arxiv.org/abs/2006.14709).
- [32] GOLUB, G. H., HEATH, M. and WAHBA, G. (1979). Generalized cross-validation as a method for choosing a good ridge parameter. *Technometrics* **21** 215–223. [MR0533250 https://doi.org/10.2307/1268518](https://doi.org/10.2307/1268518)
- [33] GOODFELLOW, I., BENGIO, Y. and COURVILLE, A. (2016). *Deep Learning. Adaptive Computation and Machine Learning*. MIT Press, Cambridge, MA. [MR3617773](https://doi.org/10.26434/chemrxiv-2016-03-001)
- [34] GUNASEKAR, S., LEE, J., SOUDRY, D. and SREBRO, N. (2018). Characterizing implicit bias in terms of optimization geometry. In *35th International Conference on Machine Learning, ICML 2018* 2932–2955. International Machine Learning Society (IMLS).
- [35] GUNASEKAR, S., LEE, J. D., SOUDRY, D. and SREBRO, N. (2018). Implicit bias of gradient descent on linear convolutional networks. In *Advances in Neural Information Processing Systems* 9461–9471.
- [36] HASTIE, T., MONTANARI, A., ROSSET, S. and TIBSHIRANI, R. J. (2022). Supplement to “Surprises in high-dimensional ridgeless least squares interpolation.” <https://doi.org/10.1214/21-AOS2133SUPP>
- [37] HASTIE, T., ROSSET, S., TIBSHIRANI, R. and ZHU, J. (2003/04). The entire regularization path for the support vector machine. *J. Mach. Learn. Res.* **5** 1391–1415. [MR2248021](https://doi.org/10.1162/jmlr.2004.5.1391)
- [38] HU, H. and LU, Y. M. (2020). Universality laws for high-dimensional learning with random features. Available at [arXiv:2009.07669](https://arxiv.org/abs/2009.07669).
- [39] JACOT, A., GABRIEL, F. and HONGLER, C. (2018). Neural tangent kernel: Convergence and generalization in neural networks. *Adv. Neural Inf. Process. Syst.* **31**.
- [40] KE, W. and THRAMOULIDIS, C. (2020). Benign overfitting in binary classification of gaussian mixtures. arXiv preprint. Available at [arXiv:2011.09148](https://arxiv.org/abs/2011.09148).
- [41] KNOWLES, A. and YIN, J. (2017). Anisotropic local laws for random matrices. *Probab. Theory Related Fields* **169** 257–352. [MR3704770 https://doi.org/10.1007/s00440-016-0730-4](https://doi.org/10.1007/s00440-016-0730-4)
- [42] KOBAC, D., LOMOND, J. and SANCHEZ, B. (2020). 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. [MR4209455](https://doi.org/10.1162/jmlr.2020.21.169)
- [43] LEDOIT, O. and PÉCHÉ, S. (2011). Eigenvectors of some large sample covariance matrix ensembles. *Probab. Theory Related Fields* **151** 233–264. [MR2834718 https://doi.org/10.1007/s00440-010-0298-3](https://doi.org/10.1007/s00440-010-0298-3)
- [44] LEE, J., XIAO, L., SCHOENHOLZ, S. S., BAHRI, Y., NOVAK, R., SOHL-DICKSTEIN, J. and PENNINGTON, J. (2020). Wide neural networks of any depth evolve as linear models under gradient descent. *J. Stat. Mech. Theory Exp.* **12** 124002. [MR4241355 https://doi.org/10.1088/1742-5468/abc62b](https://doi.org/10.1088/1742-5468/abc62b)

- [45] LI, K.-C. (1986). Asymptotic optimality of C_L and generalized cross-validation in ridge regression with application to spline smoothing. *Ann. Statist.* **14** 1101–1112. [MR0856808](#) <https://doi.org/10.1214/aos/1176350052>
- [46] LI, K.-C. (1987). Asymptotic optimality for C_p , C_L , cross-validation and generalized cross-validation: Discrete index set. *Ann. Statist.* **15** 958–975. [MR0902239](#) <https://doi.org/10.1214/aos/1176350486>
- [47] LIANG, T. and RAKHLIN, A. (2020). Just interpolate: Kernel “ridgeless” regression can generalize. *Ann. Statist.* **48** 1329–1347. [MR4124325](#) <https://doi.org/10.1214/19-AOS1849>
- [48] LIANG, T., RAKHLIN, A. and ZHAI, X. (2020). On the multiple descent of minimum-norm interpolants and restricted lower isometry of kernels. In *Conference on Learning Theory* 2683–2711.
- [49] MEI, S. and MONTANARI, A. (2019). The generalization error of random features regression: Precise asymptotics and double descent curve. *Comm. Pure Appl. Math.* To appear.
- [50] MEI, S., MONTANARI, A. and NGUYEN, P.-M. (2018). A mean field view of the landscape of two-layer neural networks. *Proc. Natl. Acad. Sci. USA* **115** E7665–E7671. [MR3845070](#) <https://doi.org/10.1073/pnas.1806579115>
- [51] MIOLANE, L. and MONTANARI, A. (2021). The distribution of the Lasso: Uniform control over sparse balls and adaptive parameter tuning. *Ann. Statist.* **49** 2313–2335. [MR4319252](#) <https://doi.org/10.1214/20-aos2038>
- [52] MONTANARI, A., RUAN, F., SOHN, Y. and YAN, J. (2019). The generalization error of max-margin linear classifiers: High-dimensional asymptotics in the overparametrized regime. Available at [arXiv:1911.01544](#).
- [53] MUTHUKUMAR, V., NARANG, A., SUBRAMANIAN, V., BELKIN, M., HSU, D. and SAHAI, A. (2021). Classification vs regression in overparameterized regimes: Does the loss function matter? *J. Mach. Learn. Res.* **22** 222. [MR4329801](#)
- [54] PATIL, P., WEI, Y., RINALDO, A. and TIBSHIRANI, R. (2021). Uniform consistency of cross-validation estimators for high-dimensional ridge regression. In *International Conference on Artificial Intelligence and Statistics* 3178–3186. PMLR.
- [55] PENNINGTON, J. and WORAH, P. (2017). Nonlinear random matrix theory for deep learning. *Adv. Neural Inf. Process. Syst.* **30**.
- [56] RAKHLIN, A. and ZHAI, X. (2019). Consistency of interpolation with Laplace kernels is a high-dimensional phenomenon. In *Conference on Learning Theory* 2595–2623. PMLR.
- [57] RICHARDS, D., MOURTADA, J. and ROSASCO, L. (2020). Asymptotics of ridge (less) regression under general source condition. Available at [arXiv:2006.06386](#).
- [58] ROSSET, S., ZHU, J. and HASTIE, T. (2003/04). Boosting as a regularized path to a maximum margin classifier. *J. Mach. Learn. Res.* **5** 941–973. [MR2248005](#)
- [59] ROTSKOFF, G. M. and VANDEN-EIJNDEN, E. (2018). Neural networks as interacting particle systems: Asymptotic convexity of the loss landscape and universal scaling of the approximation error. Available at [arXiv:1805.00915](#).
- [60] RUBIO, F. and MESTRE, X. (2011). Spectral convergence for a general class of random matrices. *Statist. Probab. Lett.* **81** 592–602. [MR2772917](#) <https://doi.org/10.1016/j.spl.2011.01.004>
- [61] SERDOBOLSKII, V. I. (2008). *Multiparametric Statistics*. Elsevier, Amsterdam. [MR2531357](#)
- [62] SIRIGNANO, J. and SPILOPOULOS, K. (2020). Mean field analysis of neural networks: A law of large numbers. *SIAM J. Appl. Math.* **80** 725–752. [MR4074020](#) <https://doi.org/10.1137/18M1192184>
- [63] SPIGLER, S., GEIGER, M., D’ASCOLI, S., SAGUN, L., BIROLI, G. and WYART, M. (2019). A jamming transition from under- to over-parametrization affects generalization in deep learning. *J. Phys. A* **52** 474001. [MR4028949](#) <https://doi.org/10.1088/1751-8121/ab4c8b>
- [64] TAO, T. (2012). *Topics in Random Matrix Theory. Graduate Studies in Mathematics* **132**. Amer. Math. Soc., Providence, RI. [MR2906465](#) <https://doi.org/10.1090/gsm/132>
- [65] TIBSHIRANI, R. J. (2015). A general framework for fast stagewise algorithms. *J. Mach. Learn. Res.* **16** 2543–2588. [MR3450516](#)
- [66] TSIGLER, A. and BARTLETT, P. L. (2020). Benign overfitting in ridge regression. Available at [arXiv:2009.14286](#).
- [67] TULINO, A. M. and VERDU, S. (2004). Random matrix theory and wireless communications. *Foundations and Trends in Communications and Information Theory* **1** 1–182.
- [68] WEBB, A. R. (1994). Functional approximation by feed-forward networks: A least-squares approach to generalization. *IEEE Trans. Neural Netw.* **5** 363–371.
- [69] WU, D. and XU, J. (2020). On the optimal weighted ℓ_2 regularization in overparameterized linear regression. *Adv. Neural Inf. Process. Syst.* **33** 10112–10123.
- [70] XU, J., MALEKI, A., RAD, K. R. and HSU, D. (2021). Consistent risk estimation in moderately high-dimensional linear regression. *IEEE Trans. Inf. Theory* **67** 5997–6030. [MR4345048](#) <https://doi.org/10.1109/TIT.2021.3095375>

- [71] ZHANG, C., BENGIO, S., HARDT, M., RECHT, B. and VINYALS, O. (2016). Understanding deep learning requires rethinking generalization. Available at [arXiv:1611.03530](https://arxiv.org/abs/1611.03530).
- [72] ZHANG, C., BENGIO, S. and SINGER, Y. (2019). Are all layers created equal? Available at [arXiv:1902.01996](https://arxiv.org/abs/1902.01996).
- [73] ZOU, D., CAO, Y., ZHOU, D. and GU, Q. (2018). Stochastic gradient descent optimizes over-parameterized deep ReLU networks. Available at [arXiv:1811.08888](https://arxiv.org/abs/1811.08888).