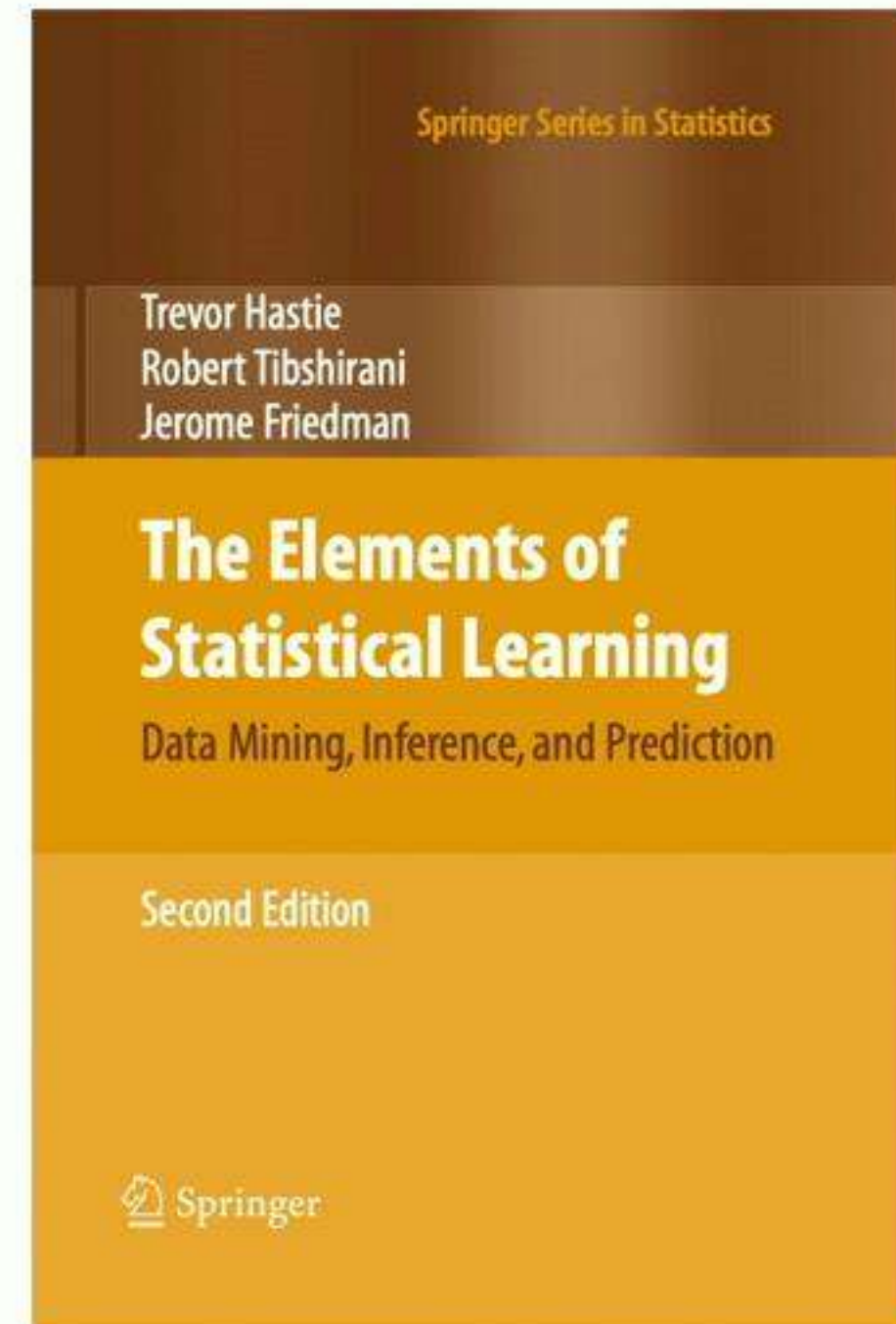


Spoilers

"A model with zero training error is overfit to the training data and will typically generalize poorly."

– Hastie, Tibshirani, & Friedman,
The Elements of Statistical Learning

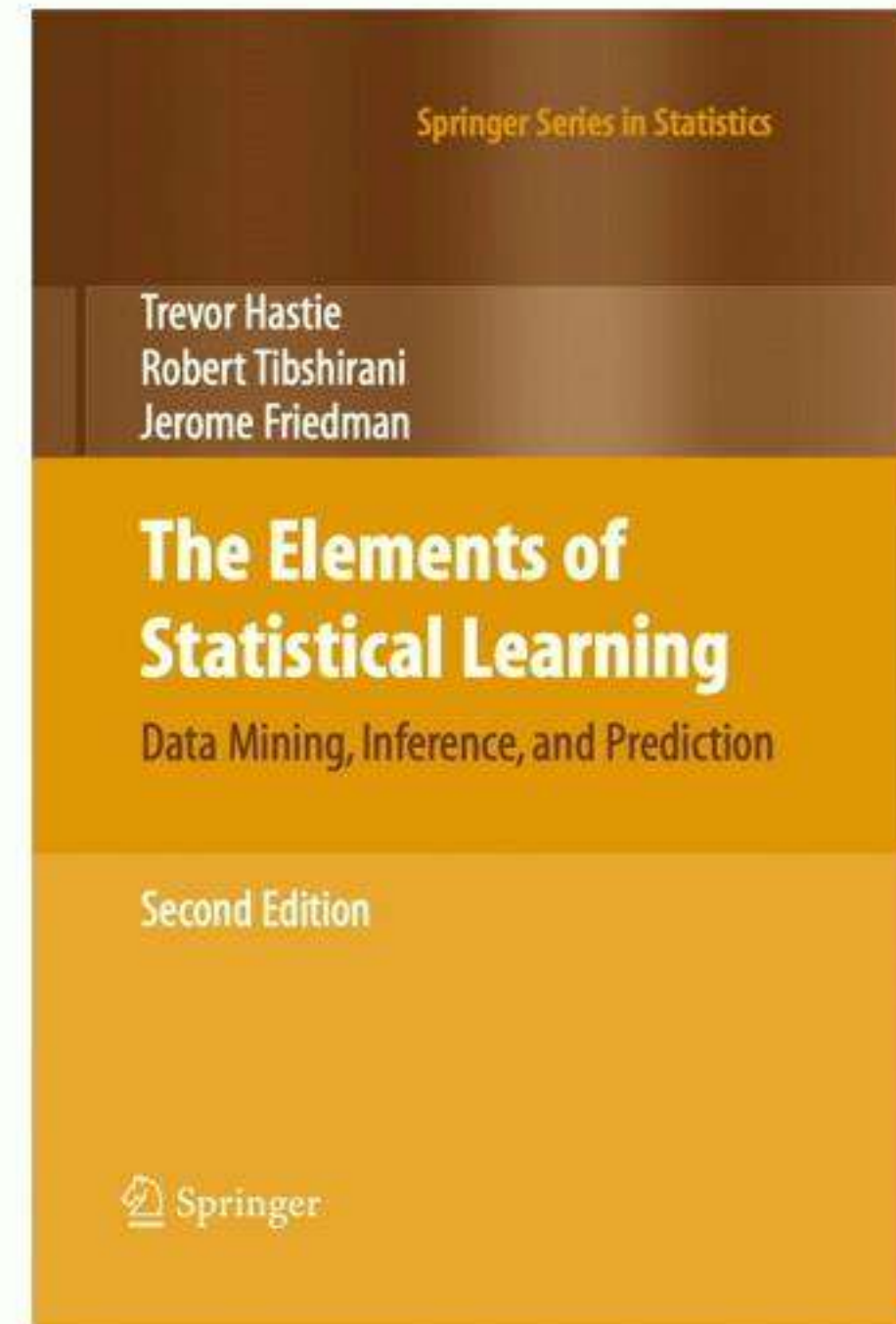


Spoilers

"A model with zero training error is overfit to the training data and will typically generalize poorly."

– Hastie, Tibshirani, & Friedman,
The Elements of Statistical Learning

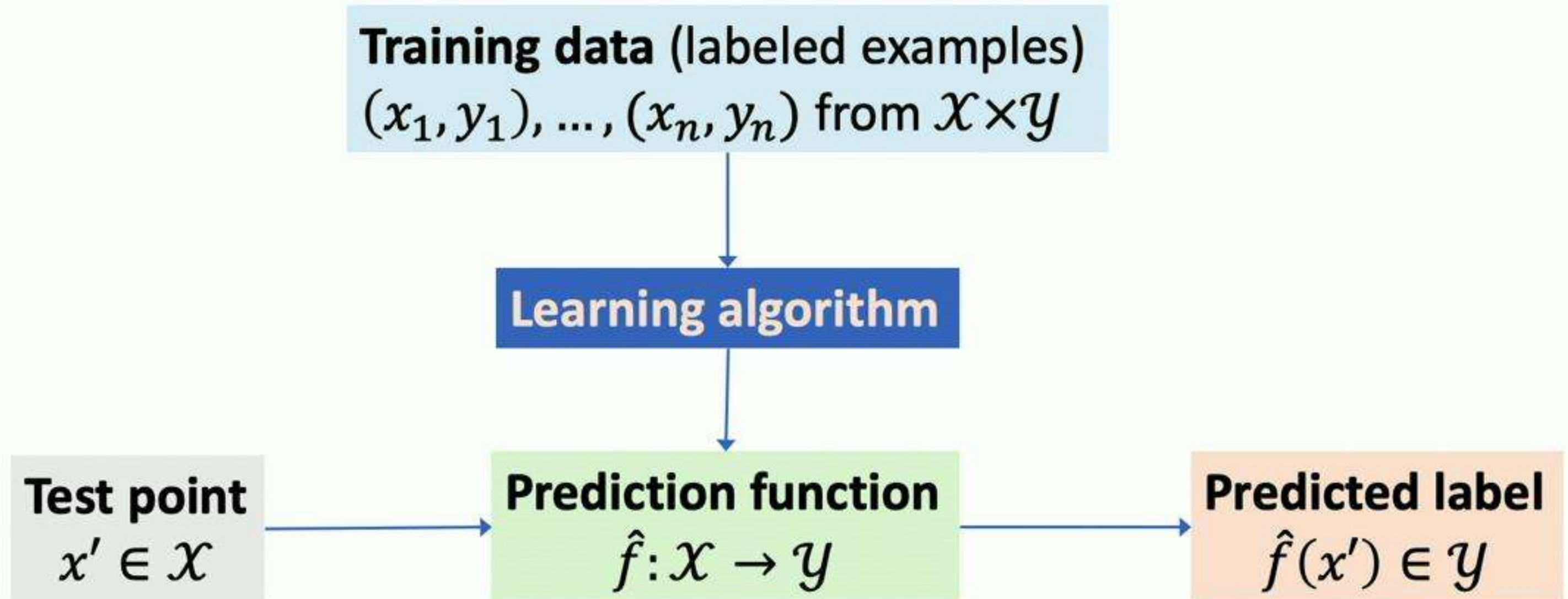
We'll give empirical and theoretical evidence against this conventional wisdom, at least in "modern" settings of machine learning.



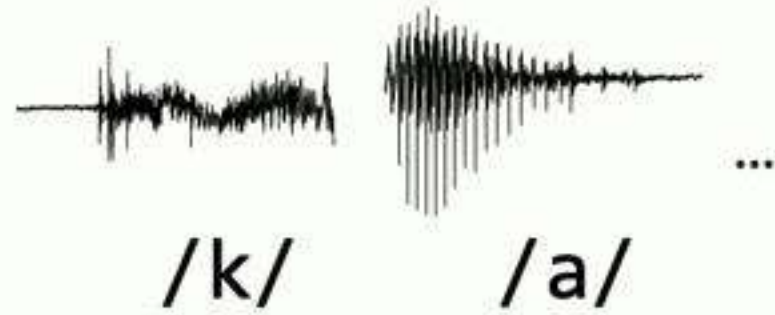
Outline

1. Observations that counter the conventional wisdom
2. Interpolation via local prediction
3. Interpolation via neural nets

Supervised learning



Supervised learning



Training data (labeled examples)
 $(x_1, y_1), \dots, (x_n, y_n)$ from $\mathcal{X} \times \mathcal{Y}$

(IID from P)

$$w \leftarrow w - \eta \nabla \hat{\mathcal{R}}(w)$$

Learning algorithm

*Risk: $\mathcal{R}(f) := \mathbb{E}[\ell(f(x'), y')]$
where $(x', y') \sim P$*

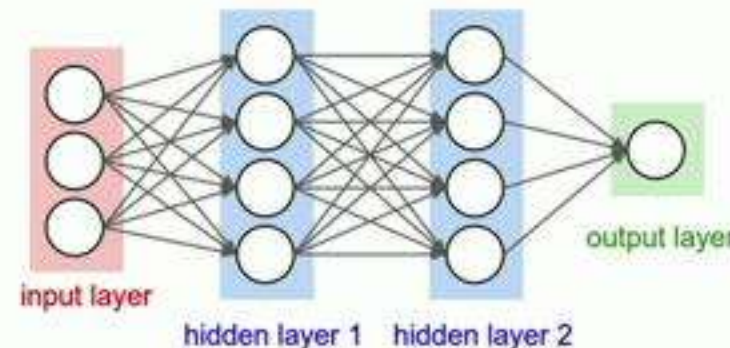
Test point
 $x' \in \mathcal{X}$

Prediction function
 $\hat{f}: \mathcal{X} \rightarrow \mathcal{Y}$

Predicted label
 $\hat{f}(x') \in \mathcal{Y}$



/t/

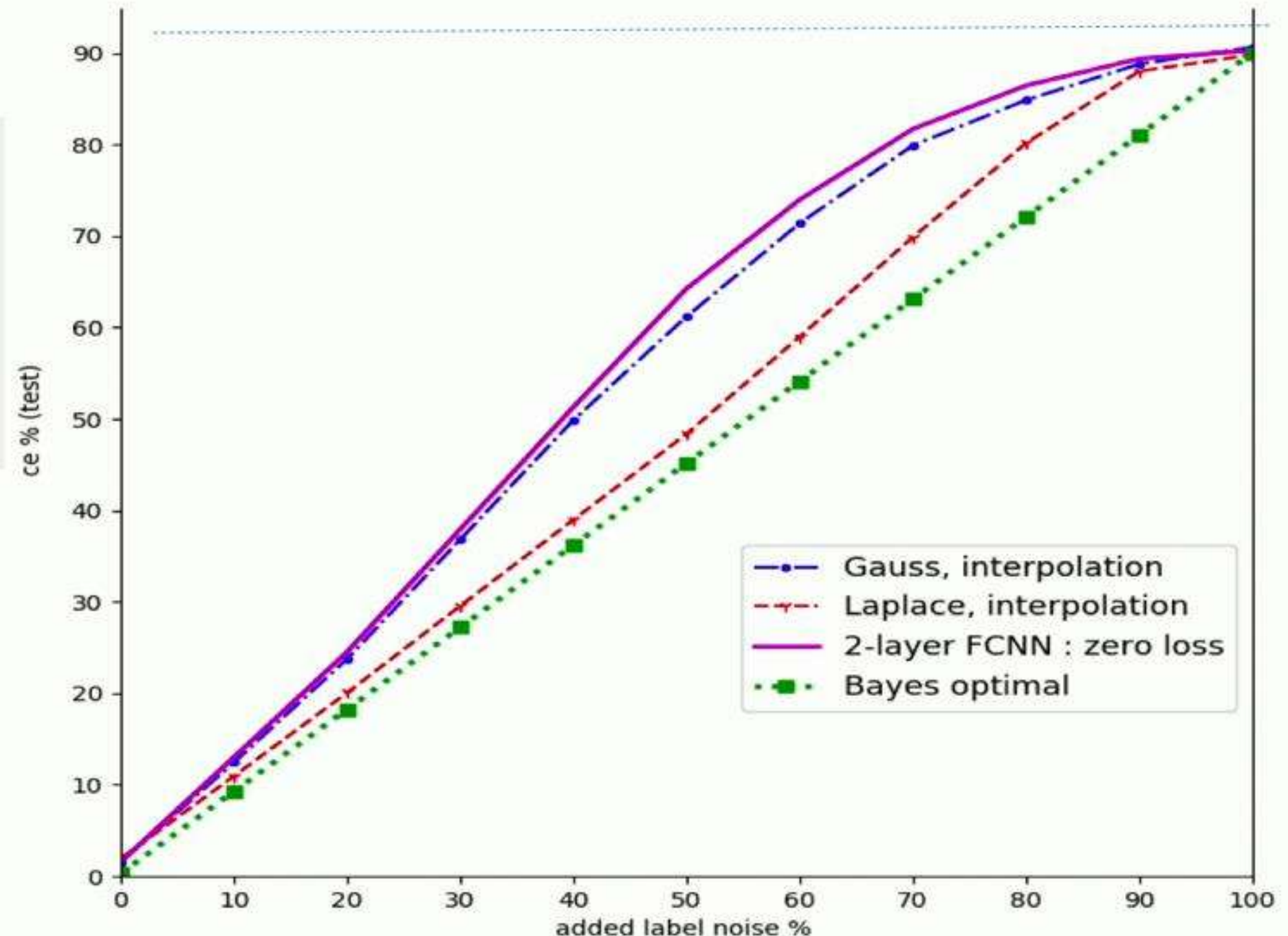


Observations from the field

(Zhang, Bengio, Hardt, Recht, & Vinyals, 2017;
Belkin, Ma, & Mandal, 2018)

Neural nets & kernel machines:

- Can **interpolate** any training data.
- Can **generalize** even when **training data has substantial amount of label noise**.

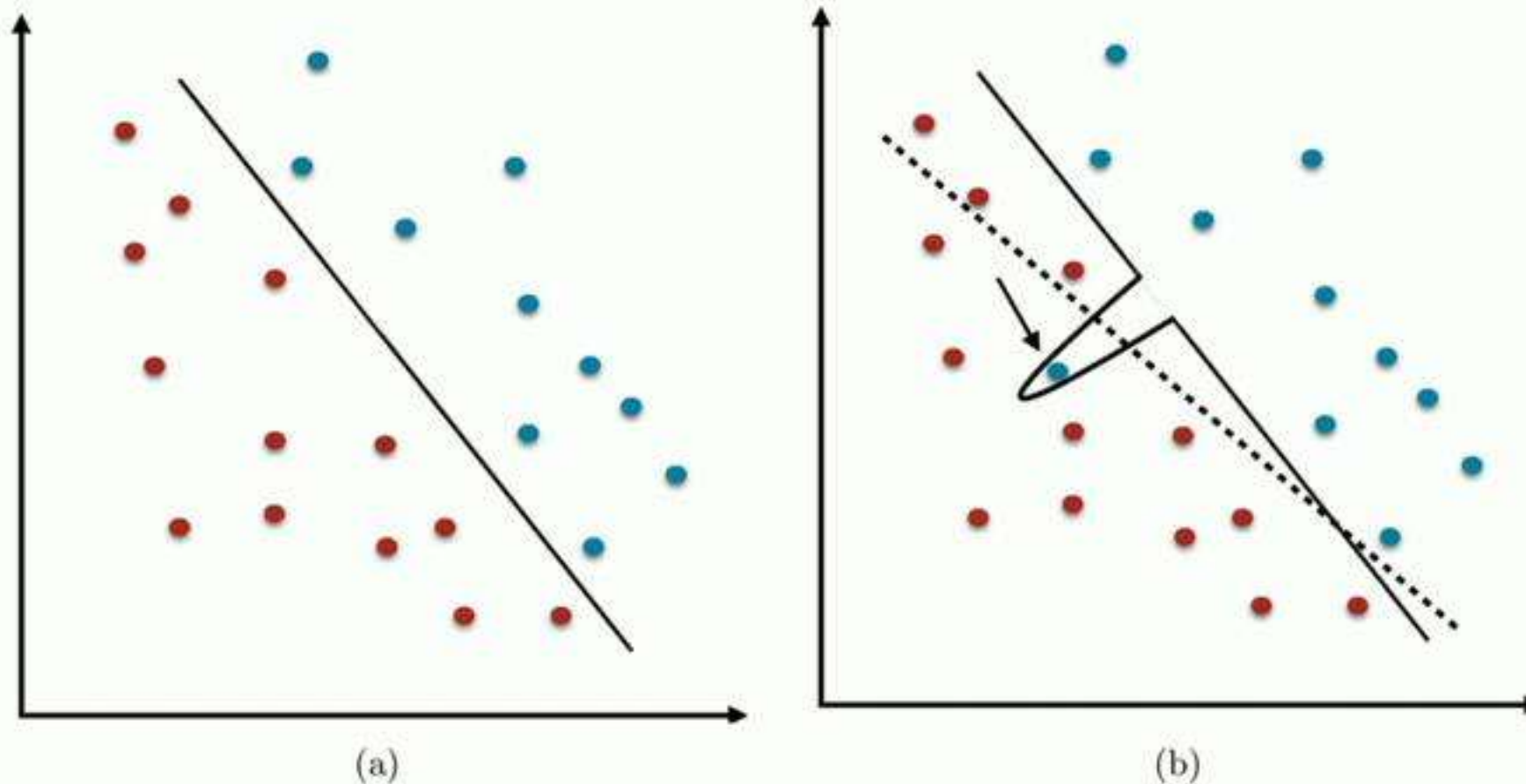


More observations from the field

(Wyner, Olson, Bleich, & Mease, 2017)

AdaBoost + large decision trees / Random forests:

- Interpret as **local interpolation** methods
- Flexibility -> **robustness to label noise**



Existing theory about local interpolation

Nearest neighbor (Cover & Hart, 1967)

- Predict with label of nearest training example
- Interpolates training data
- Risk $\rightarrow 2 \cdot \mathcal{R}(g^*)$ (sort of)



Hilbert kernel (Devroye, Györfi, & Krzyżak, 1998)

- Special kind of smoothing kernel regression (like Shepard's method)
- Interpolates training data
- Consistent*, but no convergence rates

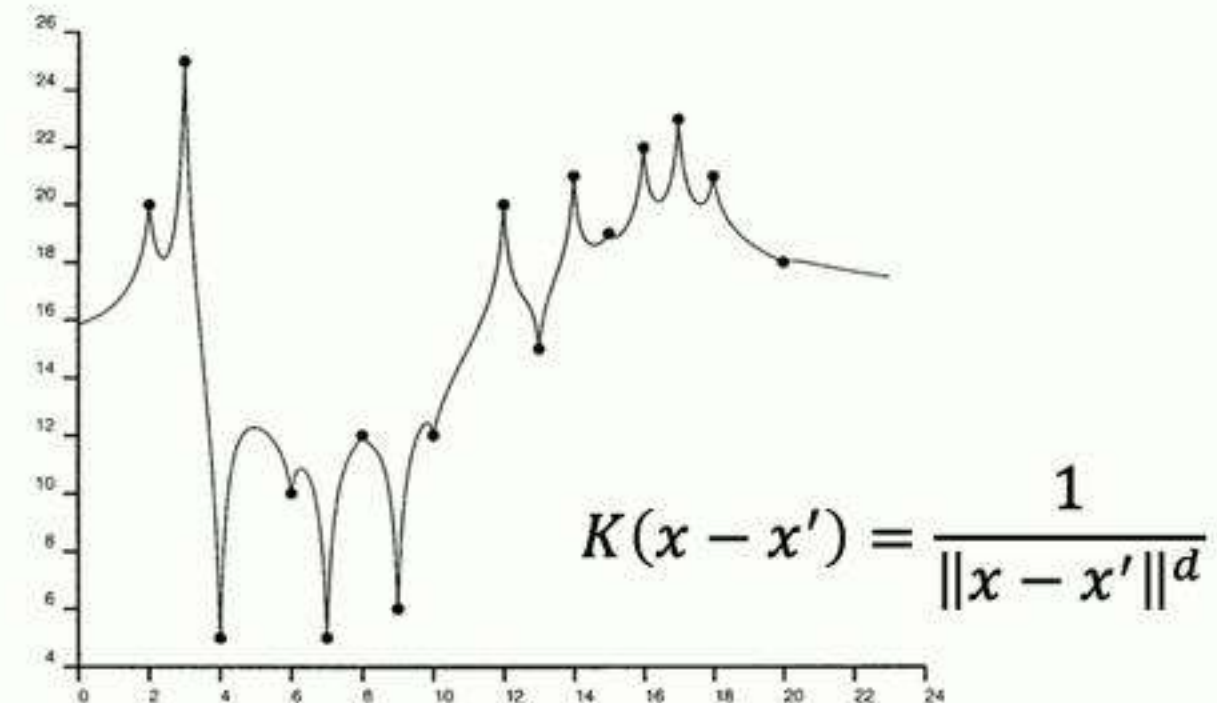


FIG. 2. The Hilbert kernel regression estimate with $d = 1$.

Our goals

- **Further counter the "conventional wisdom" re: interpolation**
Show interpolation methods can be consistent (or almost consistent) for classification & regression
 - Simplicial interpolation
 - Weighted & interpolated nearest neighbor
 - Neural nets / linear models

Our goals

- **Further counter the "conventional wisdom" re: interpolation**
Show interpolation methods can be consistent (or almost consistent) for classification & regression
 - Simplicial interpolation
 - Weighted & interpolated nearest neighbor
 - Neural nets / linear models
- Identify some **useful properties of good interpolation** methods
- Suggest **connections to practical methods**

Interpolation via local prediction

Preliminaries

- Construct estimate $\hat{\eta}$ of the **regression function**

$$\eta(x) = \mathbb{E}[y' \mid x' = x]$$

- For **binary classification** $\mathcal{Y} = \{0,1\}$:

- $\eta(x) = \Pr(y' = 1 \mid x' = x)$

- **Optimal classifier:** $f^*(x) = \mathbb{I}_{\eta(x) > \frac{1}{2}}$

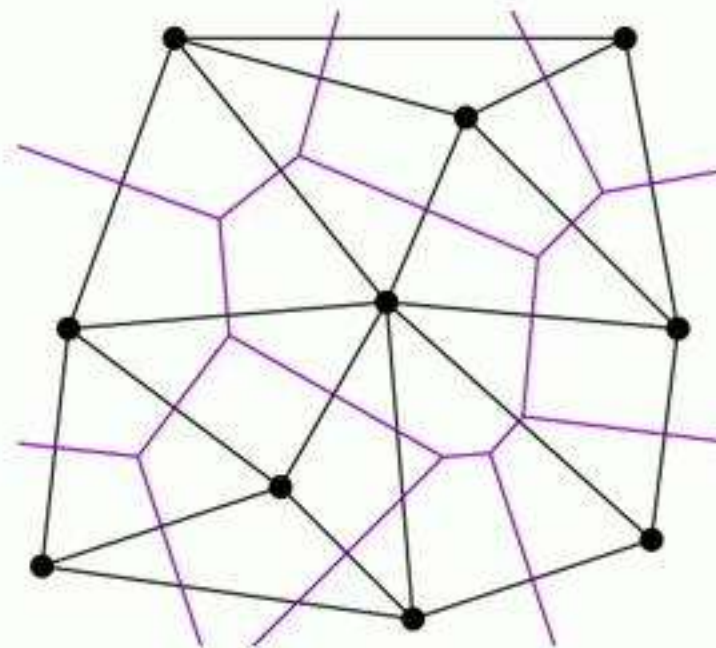
- **Plug-in classifier:** $\hat{f}(x) = \mathbb{I}_{\hat{\eta}(x) > \frac{1}{2}}$ based on estimate $\hat{\eta}$

- **Questions:**

What is the risk as $n \rightarrow \infty$? At what rate does it converge?

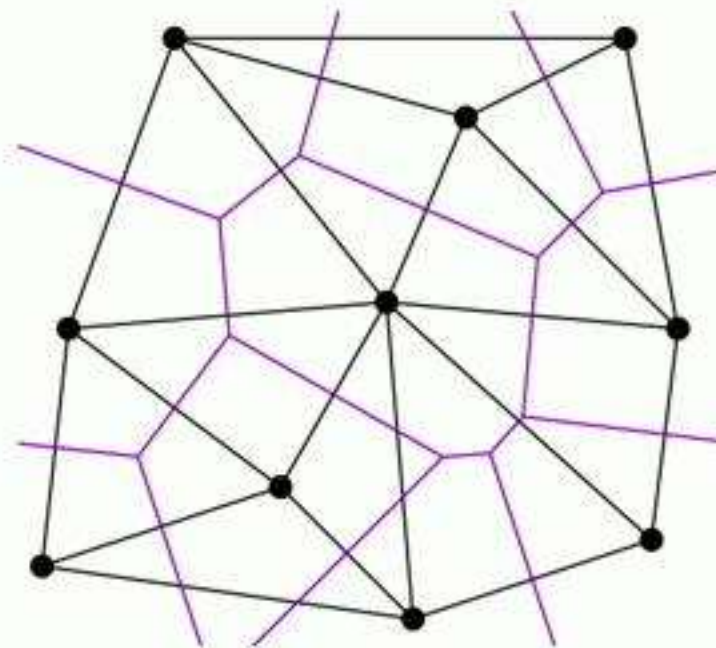
I. Simplicial interpolation

- IID training examples $(x_1, y_1), \dots, (x_n, y_n) \in \mathbb{R}^d \times [0, 1]$
 - Partition $\mathcal{C} := \text{conv}(x_1, \dots, x_n)$ into simplices with x_i as vertices via Delaunay.
 - Define $\hat{\eta}(x)$ on each simplex by affine interpolation of vertices' labels.
 - Result is piecewise linear on $\mathcal{C}$. (Punt on what happens outside of $\mathcal{C}$.)



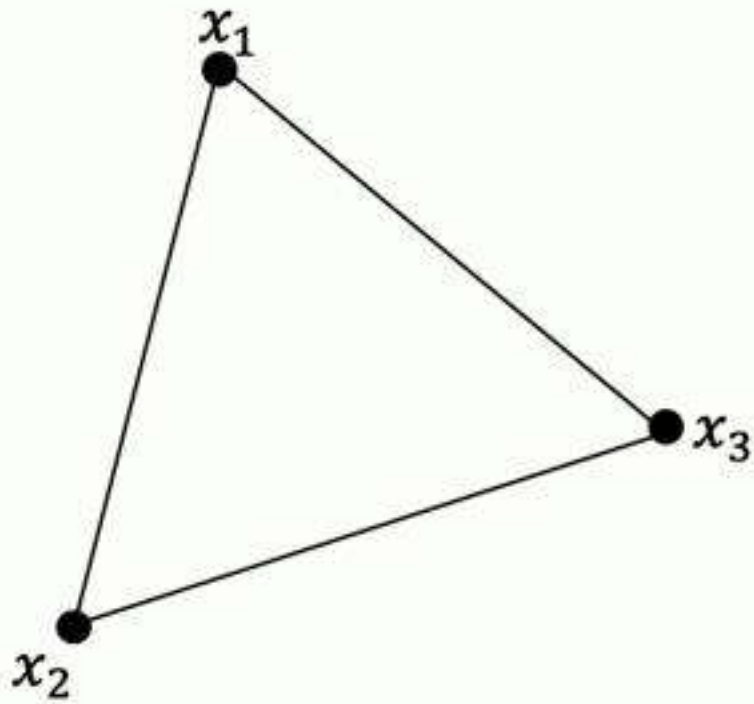
I. Simplicial interpolation

- IID training examples $(x_1, y_1), \dots, (x_n, y_n) \in \mathbb{R}^d \times [0, 1]$
 - Partition $\mathcal{C} := \text{conv}(x_1, \dots, x_n)$ into simplices with x_i as vertices via Delaunay.
 - Define $\hat{\eta}(x)$ on each simplex by affine interpolation of vertices' labels.
 - Result is piecewise linear on $\mathcal{C}$. (Punt on what happens outside of $\mathcal{C}$.)
- For classification ($y \in \{0, 1\}$), $\hat{f}$ is plug-in classifier based on $\hat{\eta}$.



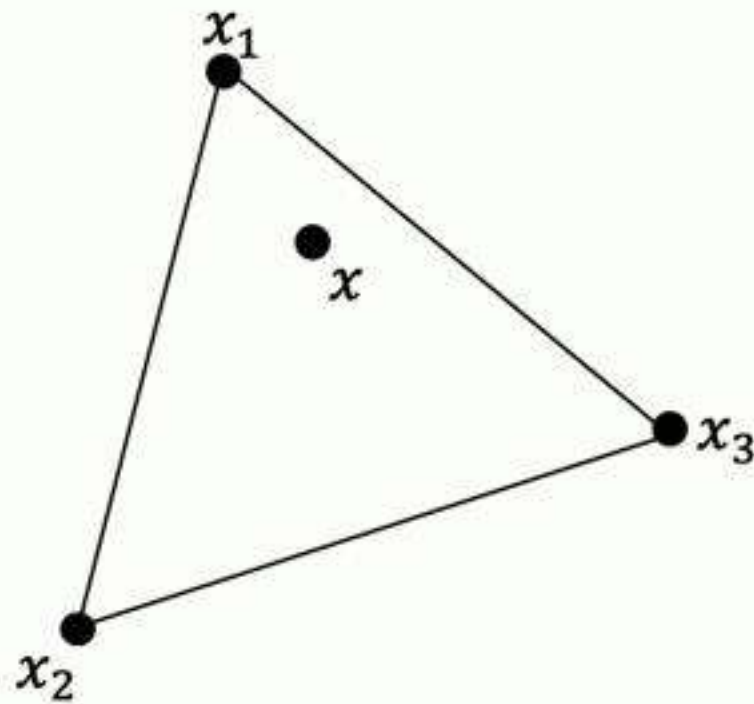
What happens on a single simplex

- Simplex on $x_1, \dots, x_{d+1}$ with corresponding labels $y_1, \dots, y_{d+1}$



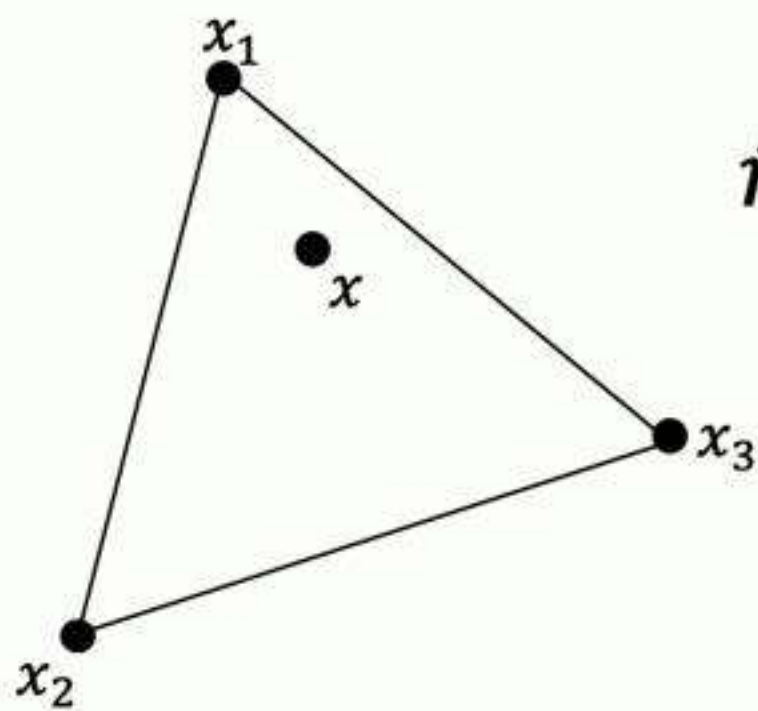
What happens on a single simplex

- Simplex on $x_1, \dots, x_{d+1}$ with corresponding labels $y_1, \dots, y_{d+1}$
- Test point x in simplex, with barycentric coordinates $(w_1, \dots, w_{d+1})$.



What happens on a single simplex

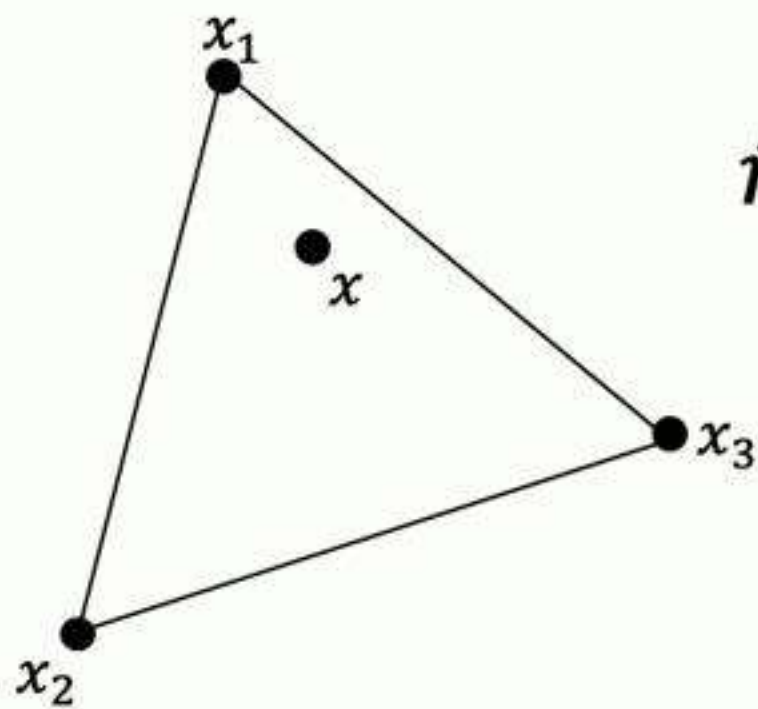
- Simplex on $x_1, \dots, x_{d+1}$ with corresponding labels $y_1, \dots, y_{d+1}$
- Test point x in simplex, with barycentric coordinates $(w_1, \dots, w_{d+1})$.
- Linear interpolation at x (i.e., least squares fit, evaluated at x):



$$\hat{\eta}(x) = \sum_{i=1}^{d+1} w_i y_i$$

What happens on a single simplex

- Simplex on $x_1, \dots, x_{d+1}$ with corresponding labels $y_1, \dots, y_{d+1}$
- Test point x in simplex, with barycentric coordinates $(w_1, \dots, w_{d+1})$.
- Linear interpolation at x (i.e., least squares fit, evaluated at x):



$$\hat{\eta}(x) = \sum_{i=1}^{d+1} w_i y_i$$

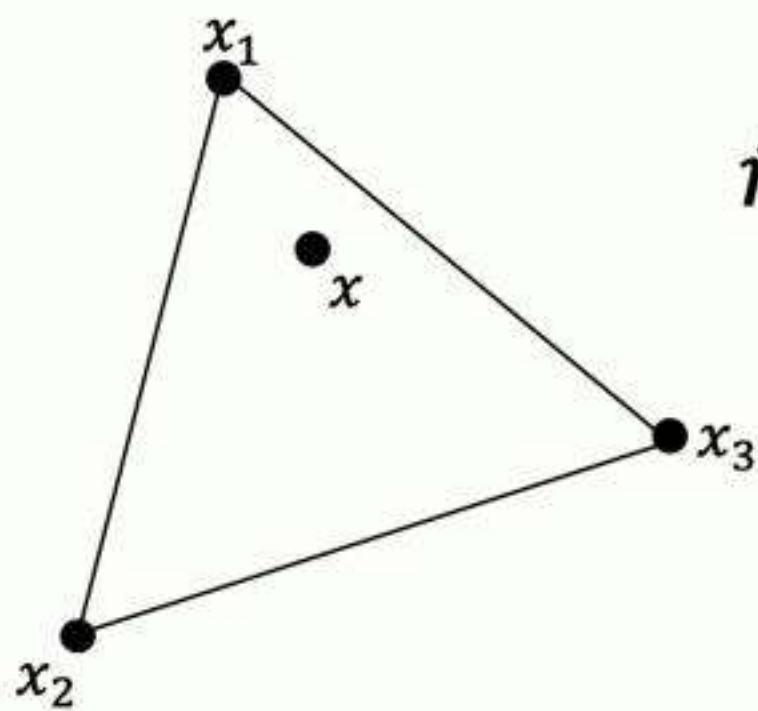
Key idea: aggregates information from all vertices to make prediction. (C.f. nearest neighbor rule.)

Comparison to nearest neighbor rule

- Suppose $\eta(x) = \Pr(y = 1 \mid x) < 1/2$ for all points in a simplex
 - Optimal prediction of f^* is 0 for all points in simplex.

What happens on a single simplex

- Simplex on $x_1, \dots, x_{d+1}$ with corresponding labels $y_1, \dots, y_{d+1}$
- Test point x in simplex, with barycentric coordinates $(w_1, \dots, w_{d+1})$.
- Linear interpolation at x (i.e., least squares fit, evaluated at x):



$$\hat{\eta}(x) = \sum_{i=1}^{d+1} w_i y_i$$

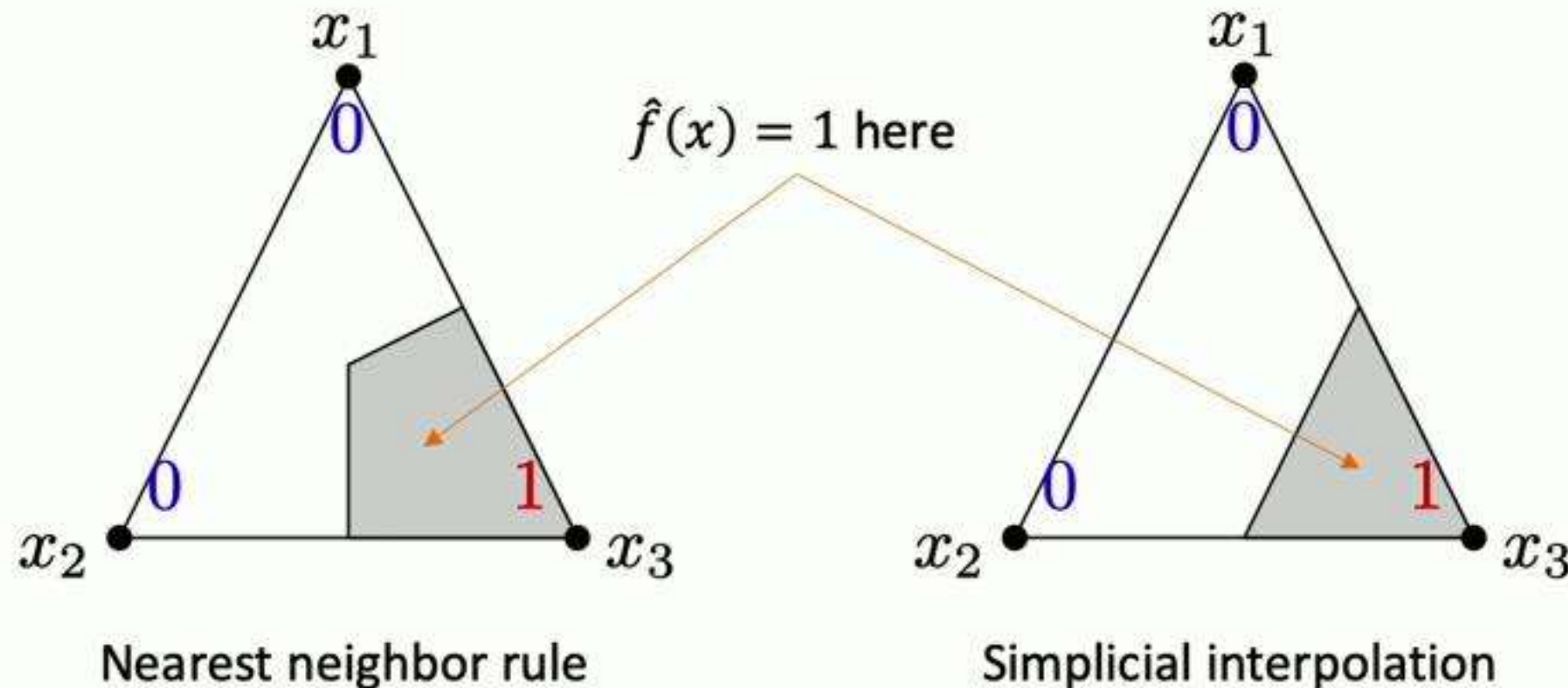
Key idea: aggregates information from all vertices to make prediction. (C.f. nearest neighbor rule.)

Comparison to nearest neighbor rule

- Suppose $\eta(x) = \Pr(y = 1 \mid x) < 1/2$ for all points in a simplex
 - Optimal prediction of f^* is 0 for all points in simplex.

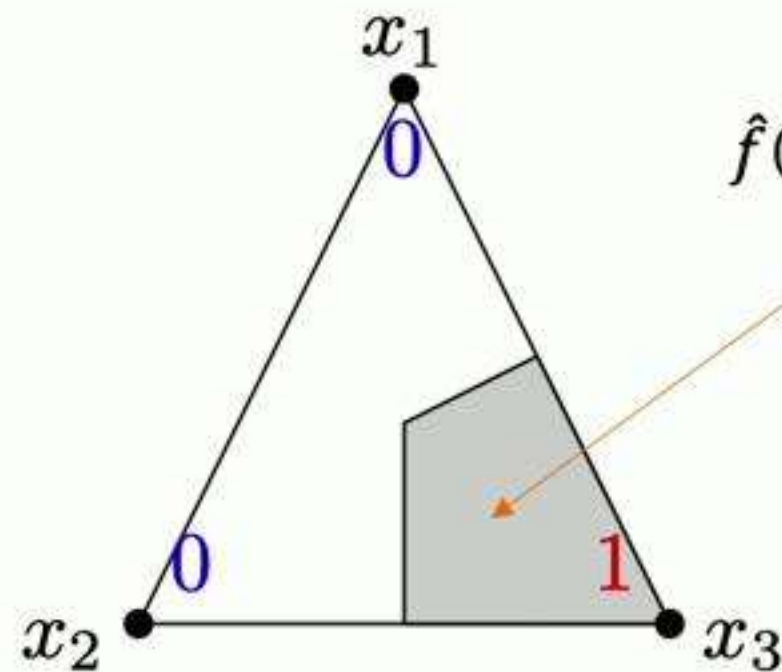
Comparison to nearest neighbor rule

- Suppose $\eta(x) = \Pr(y = 1 \mid x) < 1/2$ for all points in a simplex
 - Optimal prediction of f^* is 0 for all points in simplex.
- Suppose $y_1 = \dots = y_d = 0$, but $y_{d+1} = 1$ (due to "label noise")

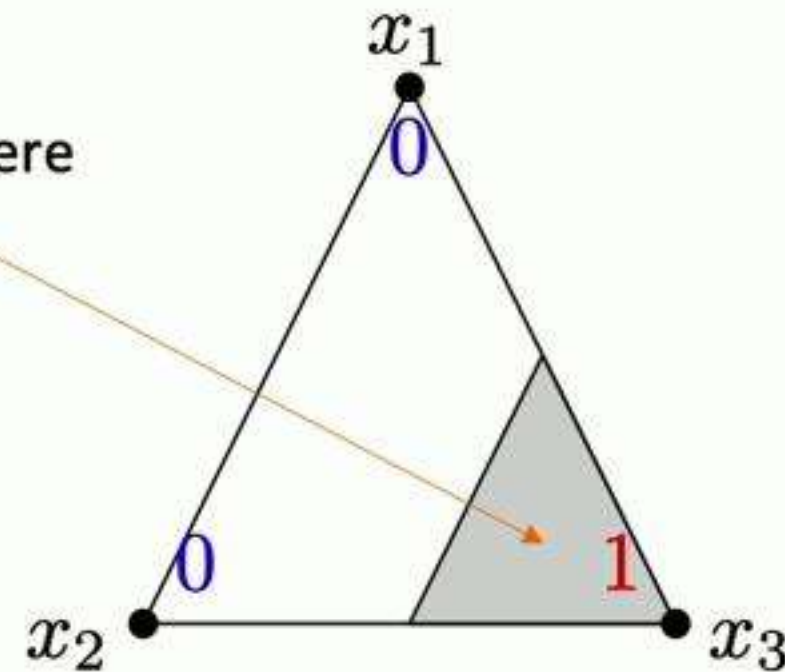


Comparison to nearest neighbor rule

- Suppose $\eta(x) = \Pr(y = 1 \mid x) < 1/2$ for all points in a simplex
 - Optimal prediction of f^* is 0 for all points in simplex.
- Suppose $y_1 = \dots = y_d = 0$, but $y_{d+1} = 1$ (due to "label noise")



Nearest neighbor rule



Simplicial interpolation

Effect is exponentially more pronounced in high dimensions!

Asymptotic risk for simplicial interpolation

[Belkin, H., Mitra, '18]

Theorem (classification): Assume distribution of x' is uniform on some convex set, and η is bounded away from $1/2$. Then simplicial interpolation's plug-in classifier $\hat{f}$ satisfies

$$\limsup_n \mathbb{E}[\text{zero/one loss}] \leq (1 + e^{-\Omega(d)}) \cdot \text{OPT}$$

Asymptotic risk for simplicial interpolation

[Belkin, H., Mitra, '18]

Theorem (classification): Assume distribution of x' is uniform on some convex set, and η is bounded away from $1/2$. Then simplicial interpolation's plug-in classifier $\hat{f}$ satisfies

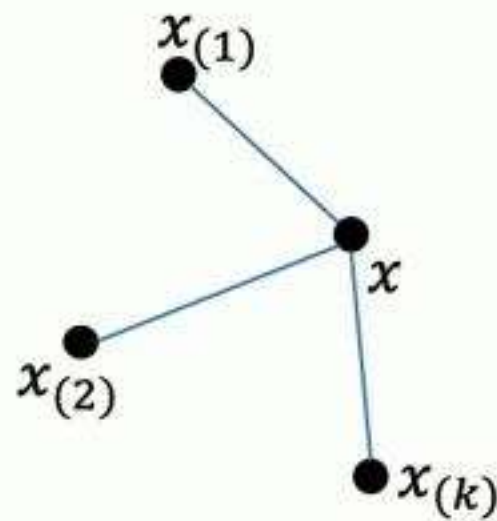
$$\limsup_n \mathbb{E}[\text{zero/one loss}] \leq (1 + e^{-\Omega(d)}) \cdot \text{OPT}$$

- **C.f. nearest neighbor classifier:** $\limsup_n \mathbb{E}[\mathcal{R}(\hat{f})] \approx 2 \cdot \mathcal{R}(f^*)$
- **For regression** (squared error):

$$\limsup_n \mathbb{E}[\text{squared error}] \leq \left(1 + o\left(\frac{1}{d}\right)\right) \cdot \text{OPT}$$

II. Weighted & interpolated NN scheme

- For given test point x , let $x_{(1)}, \dots, x_{(k)}$ be k nearest neighbors in training data, and let $y_{(1)}, \dots, y_{(k)}$ be corresponding labels.



Define

$$\hat{\eta}(x) = \frac{\sum_{i=1}^k w(x, x_{(i)}) y_{(i)}}{\sum_{i=1}^k w(x, x_{(i)})}$$

where

$$w(x, x_{(i)}) = \|x - x_{(i)}\|^{-\delta}, \quad \delta > 0$$

Interpolation: $\hat{\eta}(x) \rightarrow y_i$ as $x \rightarrow x_i$

Comparison to Hilbert kernel estimate

Weighted & interpolated NN

$$\hat{\eta}(x) = \frac{\sum_{i=1}^k w(x, x_{(i)}) y_{(i)}}{\sum_{i=1}^k w(x, x_{(i)})}$$

$$w(x, x_{(i)}) = \|x - x_{(i)}\|^{-\delta}$$

Optimal non-parametric rates

Hilbert kernel (Devroye, Györfi, & Krzyżak, 1998)

$$\hat{\eta}(x) = \frac{\sum_{i=1}^n w(x, x_i) y_i}{\sum_{i=1}^n w(x, x_i)}$$

$$w(x, x_i) = \|x - x_i\|^{-\delta}$$

Consistent, but no non-asymptotic rates

Comparison to Hilbert kernel estimate

Weighted & interpolated NN

$$\hat{\eta}(x) = \frac{\sum_{i=1}^k w(x, x_{(i)}) y_{(i)}}{\sum_{i=1}^k w(x, x_{(i)})}$$

$$w(x, x_{(i)}) = \|x - x_{(i)}\|^{-\delta}$$

Optimal non-parametric rates

Hilbert kernel (Devroye, Györfi, & Krzyżak, 1998)

$$\hat{\eta}(x) = \frac{\sum_{i=1}^n w(x, x_i) y_i}{\sum_{i=1}^n w(x, x_i)}$$

$$w(x, x_i) = \|x - x_i\|^{-\delta}$$

Consistent, but no non-asymptotic rates

Localization is essential to get non-asymptotic rate.

Convergence rates for WiNN

[Belkin, H., Mitra, '18]

Theorem: Assume distribution of x' is uniform on some compact set satisfying regularity condition, and η is α -Holder smooth.

For appropriate setting of k , weighted & interpolated NN estimate $\hat{\eta}$ satisfies

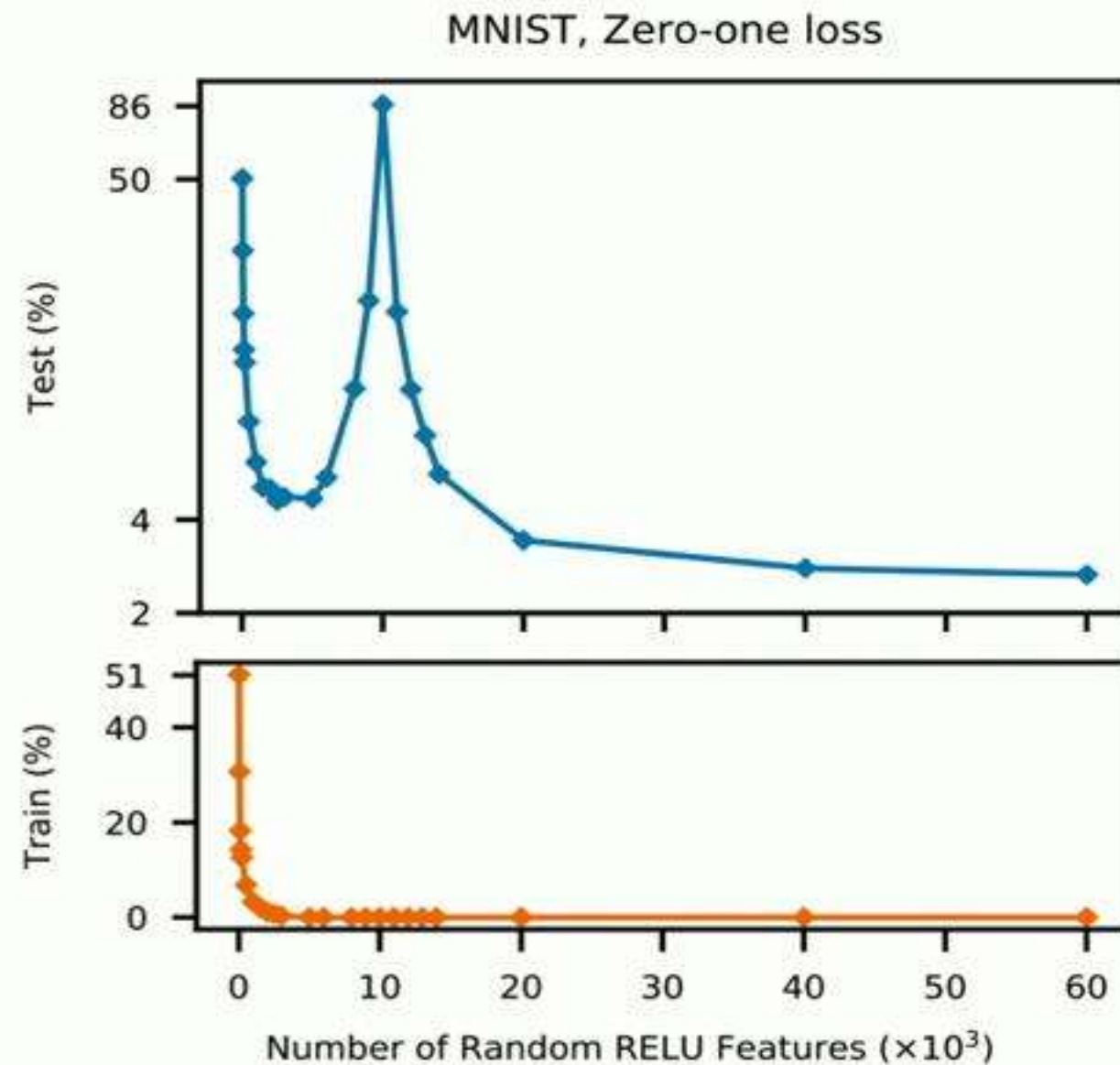
$$\mathbb{E} \left[\left(\hat{\eta}(X) - \eta(X) \right)^2 \right] \leq O\left(n^{-2\alpha/(2\alpha+d)}\right)$$

- Consistency + **optimal rates of convergence** for interpolating method.
- Follow-up work by Belkin, Rakhlin, Tsybakov '19: also for Nadaraya-Watson with **compact & singular** kernel.

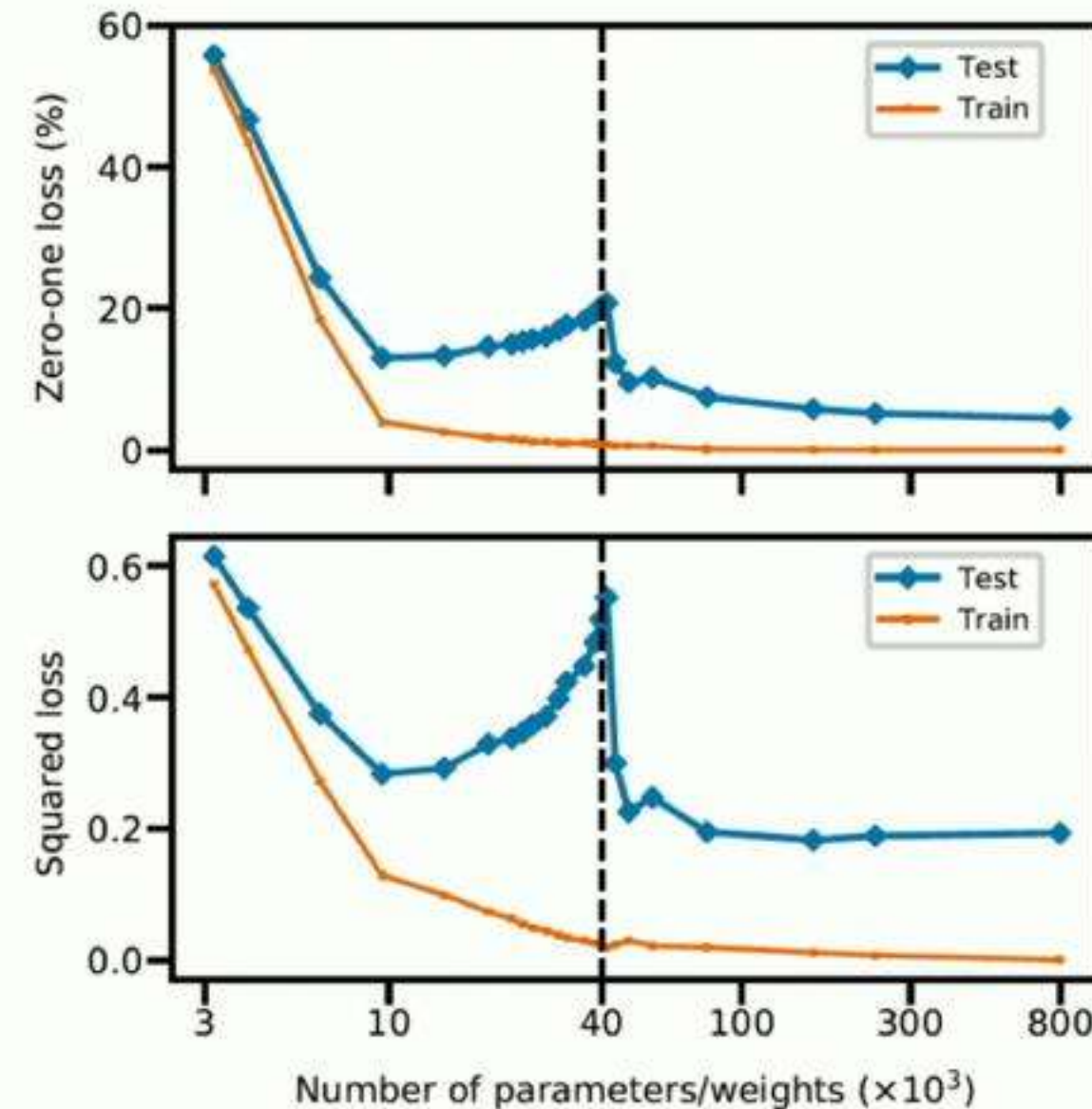
Interpolation via neural nets

Two layer fully-connected neural networks

[Belkin, H., Ma, Mandal, '19]



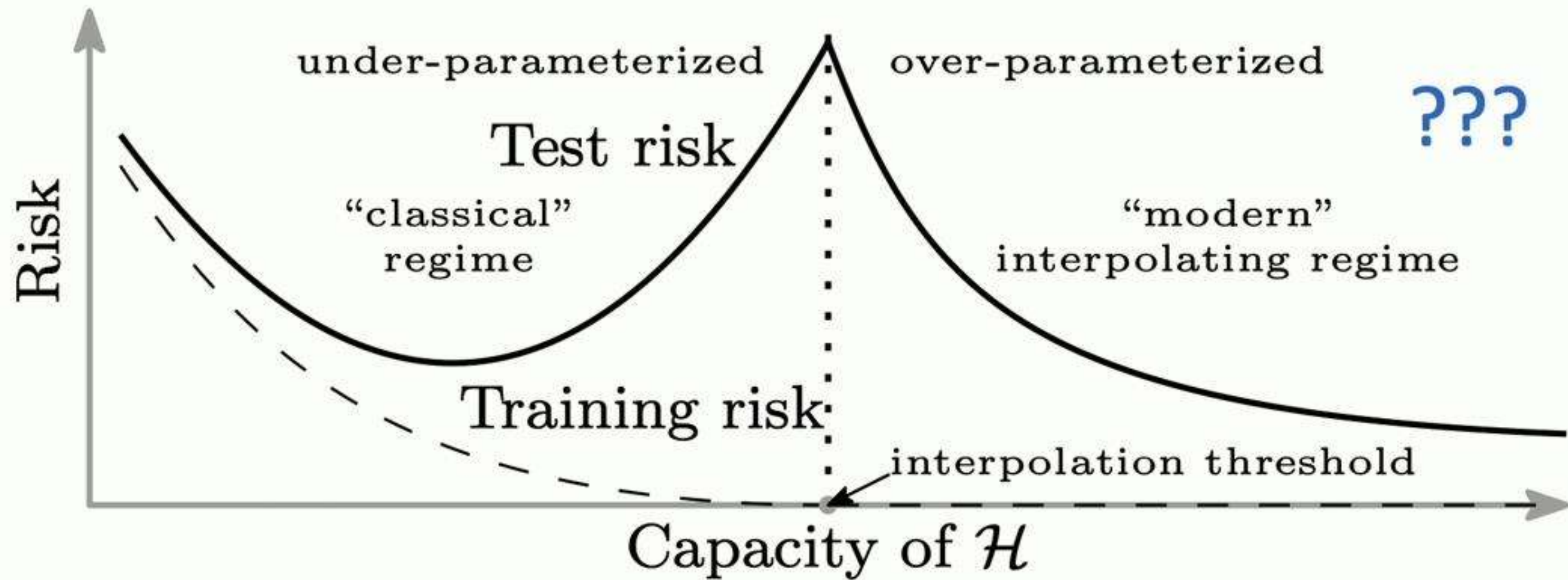
Random first layer



Trained first layer

"Double descent" risk curve

[Belkin, H., Ma, Mandal, '19]

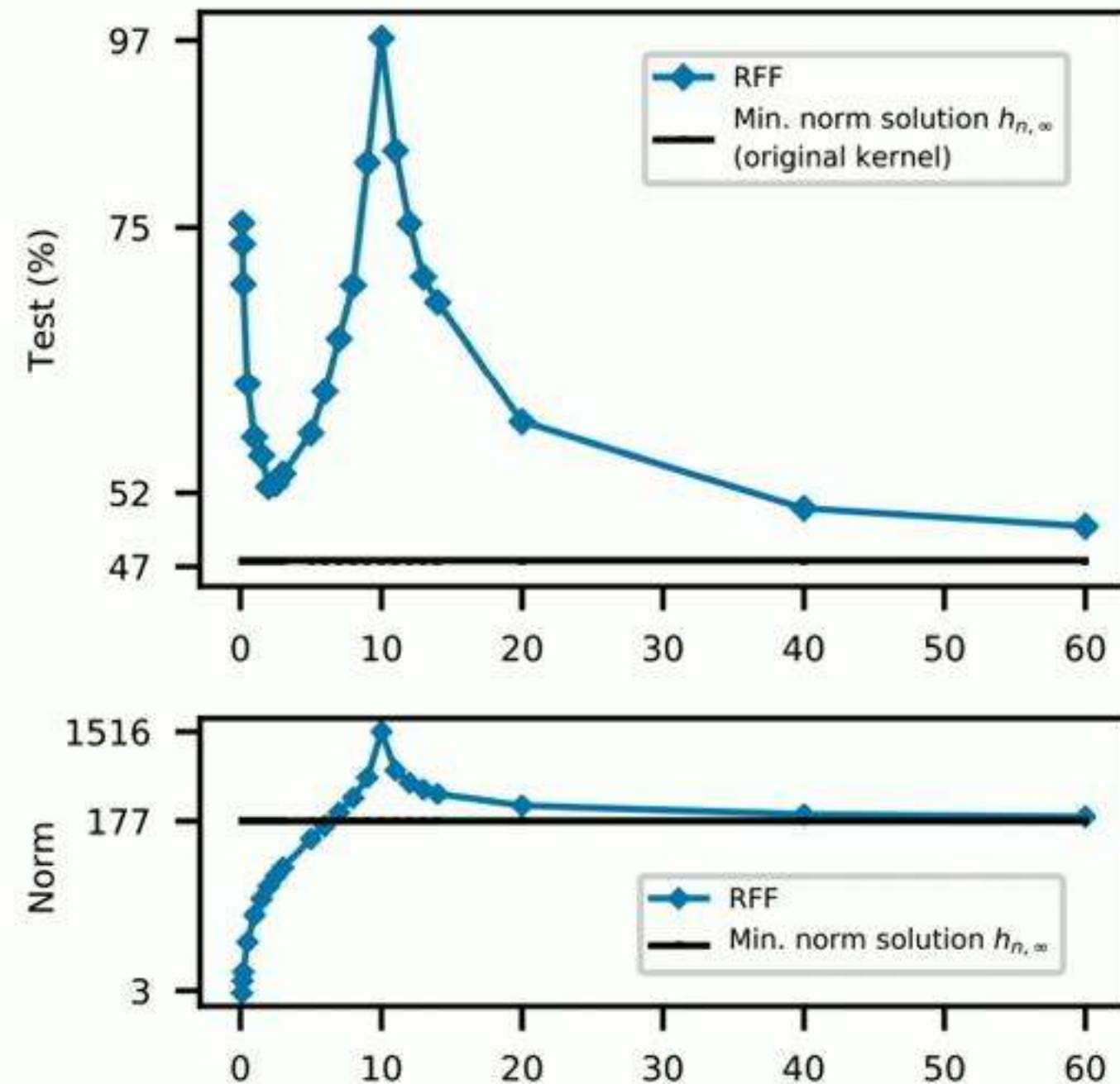


Approximating a kernel machine?

[Belkin, H., Ma, Mandal, '19]

TIMIT, Zero-one loss

- Effectiveness of interpolation depends on ability to align with the "right" inductive bias
- E.g., low RKHS norm

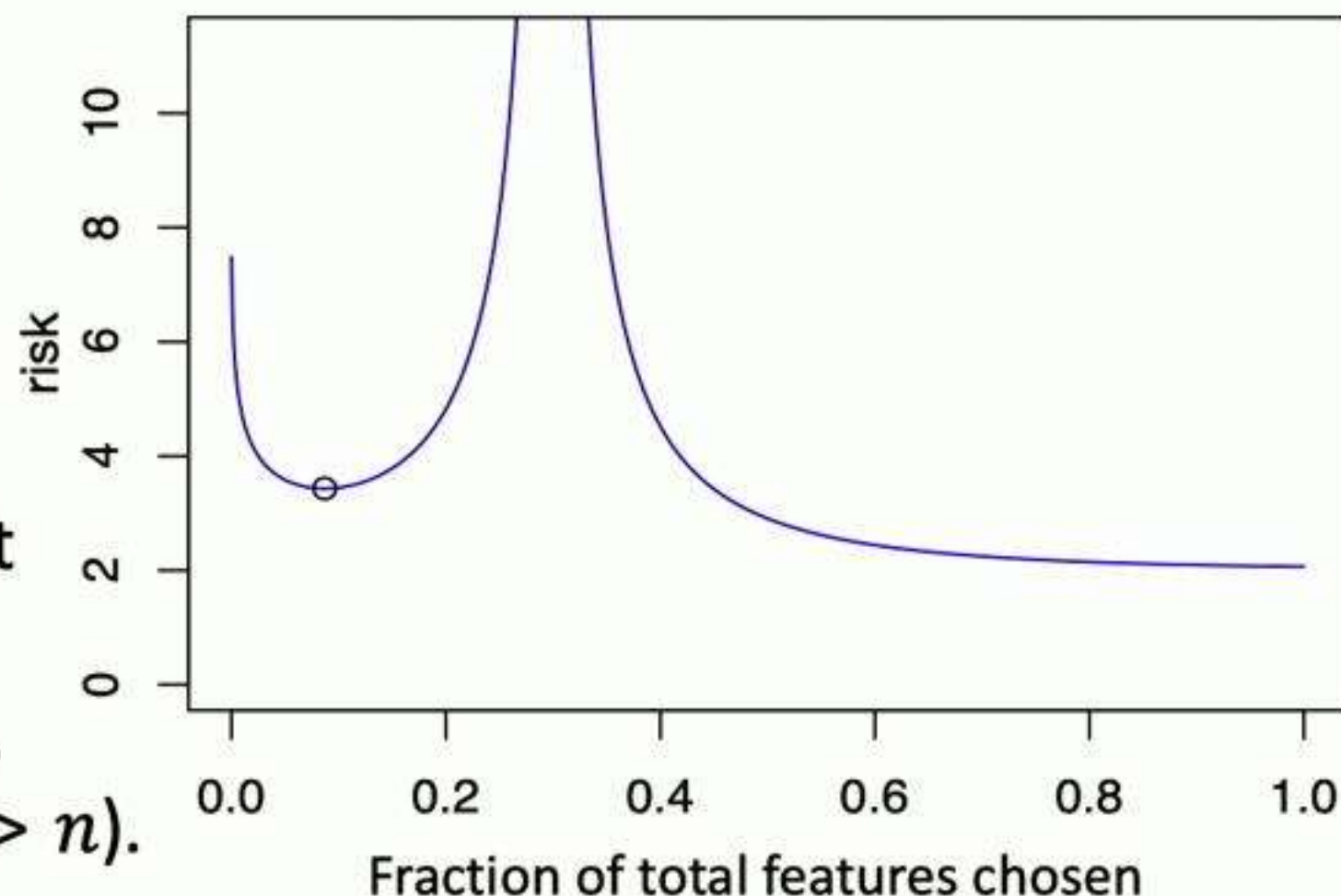


Linear regression w/ weak features [Belkin, H., Xu, '19+; Xu, H., '19+]

Gaussian design linear model with D features
All features are "relevant" but equally weak

Only use p of the features ($1 \leq p \leq D$)
Least squares ($p \leq n$) or least norm ($p \geq n$) fit

Theorem: Under certain eigenvalue condition,
minimum is beyond point of interpolation ($p > n$).



Concurrent work by Hastie, Montanari, Rosset, Tibshirani '19 (also for some non-linear models!).

Other analyses of linear models: Muthukumar, Vodrahalli, Sahai, '19; Bartlett, Long, Lugosi, Tsigler, '19.

Conclusions/open problems

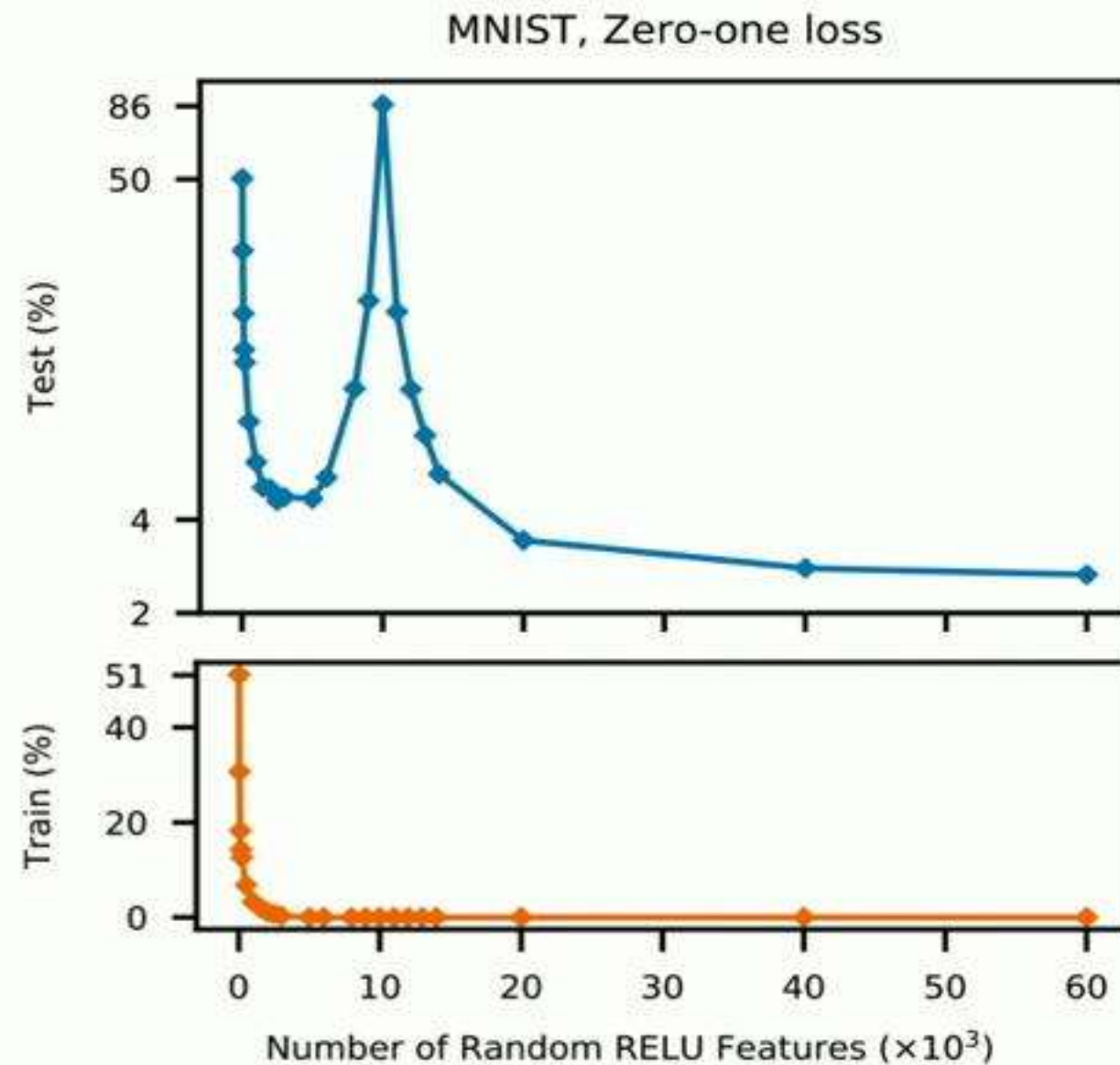
1. Interpolation is compatible with good statistical properties
2. They work by relying (exclusively!) on **inductive bias**: e.g.,
 1. Smoothness from **local averaging in high-dimensions**.
 2. Low function space norm

Open problems:

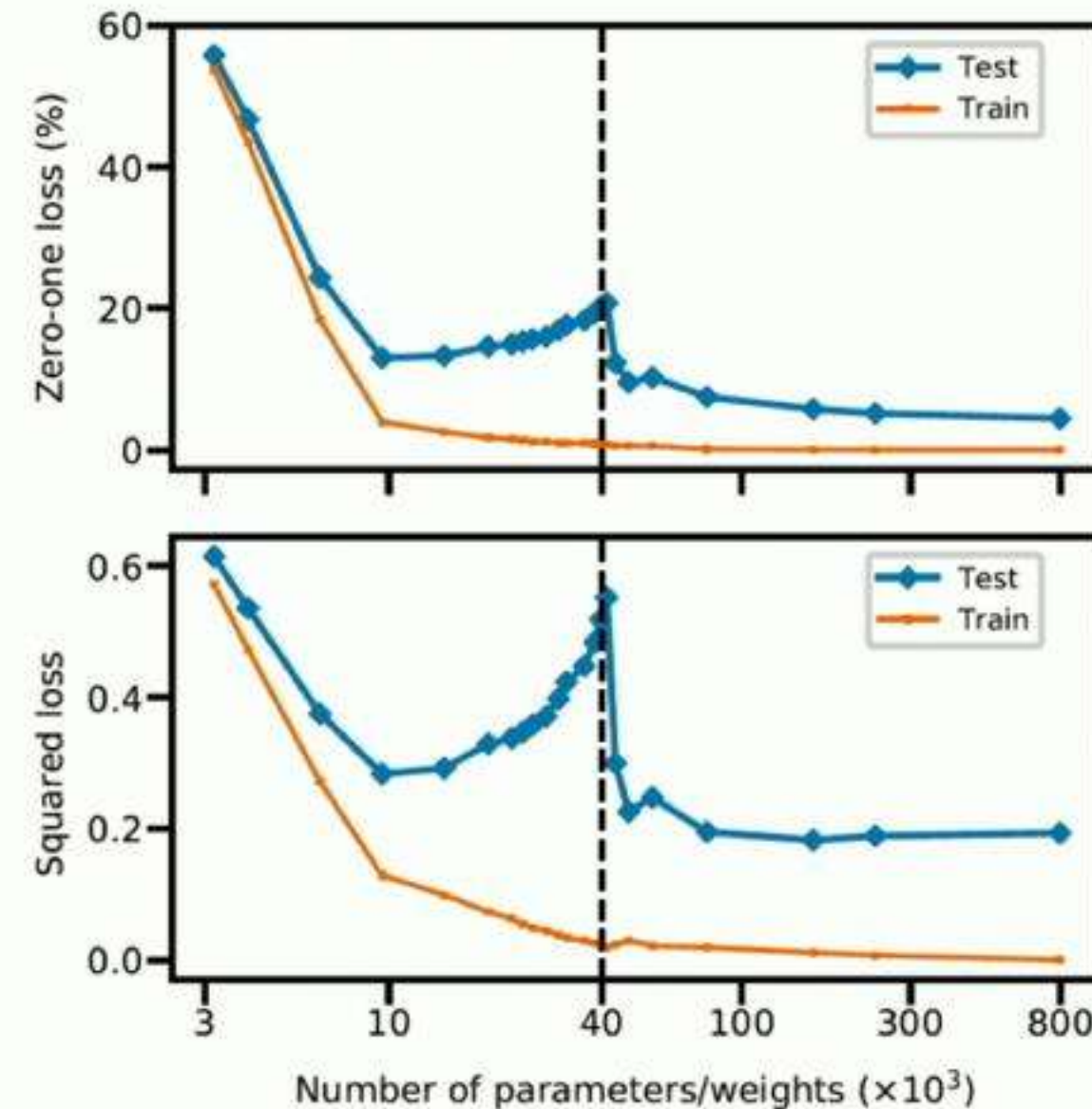
- Characterize inductive biases of other learning algorithms
- Benefits of interpolation?

Two layer fully-connected neural networks

[Belkin, H., Ma, Mandal, '19]



Random first layer



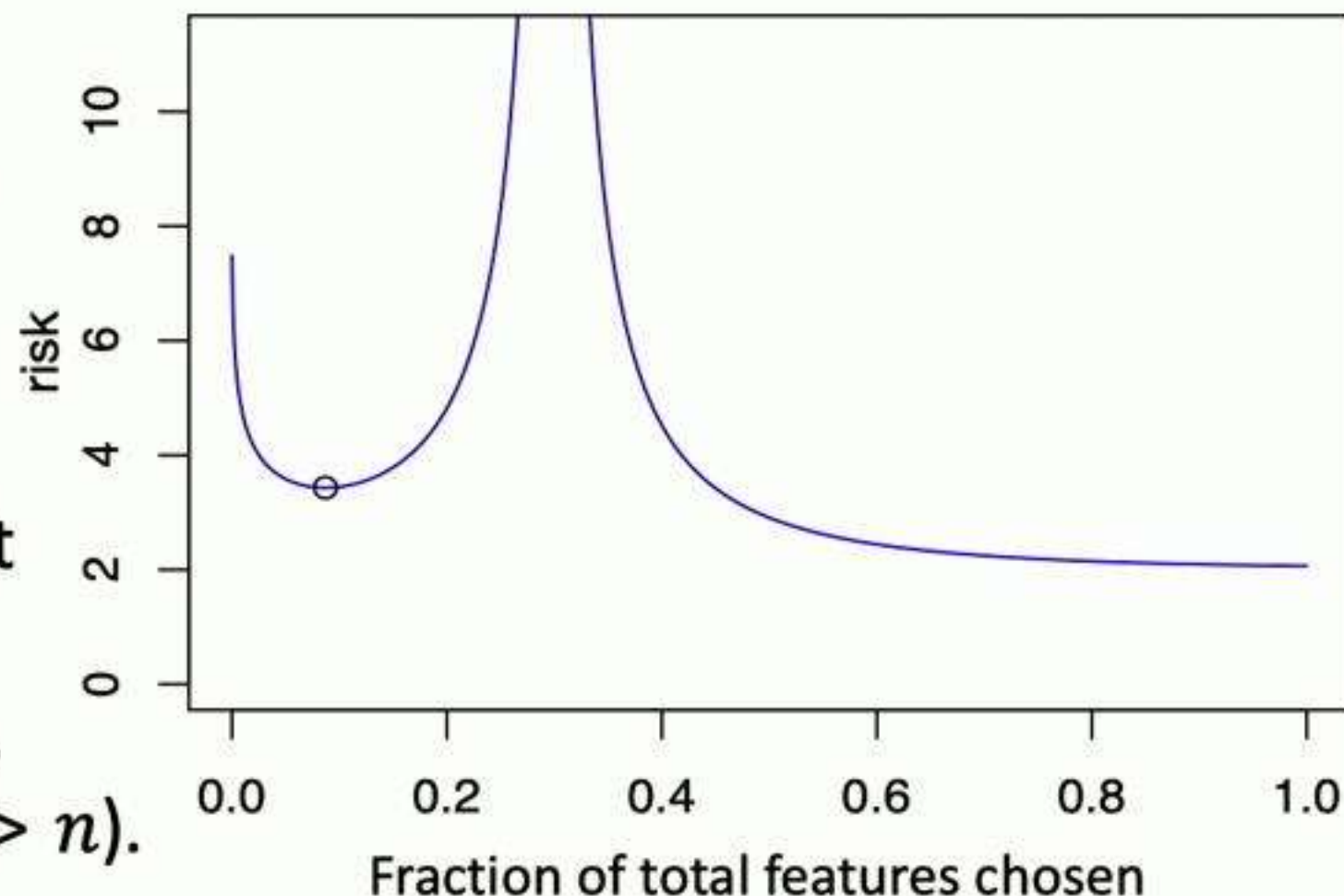
Trained first layer

Linear regression w/ weak features [Belkin, H., Xu, '19+; Xu, H., '19+]

Gaussian design linear model with D features
All features are "relevant" but equally weak

Only use p of the features ($1 \leq p \leq D$)
Least squares ($p \leq n$) or least norm ($p \geq n$) fit

Theorem: Under certain eigenvalue condition,
minimum is beyond point of interpolation ($p > n$).



Concurrent work by Hastie, Montanari, Rosset, Tibshirani '19 (also for some non-linear models!).

Other analyses of linear models: Muthukumar, Vodrahalli, Sahai, '19; Bartlett, Long, Lugosi, Tsigler, '19.

On learning methods that memorize data

Sasha Rakhlin

MIT

Aug 26, 2019

Based on (Belkin, R., Tsybakov '18), (Liang and R. '18), (R. and Zhai '18),
(Liang, R., Zhai '19)

Can a learning method be successful if it memorizes the training data?

OUTLINE

Warmup: The Nadaraya-Watson Estimator

Kernel Ridgeless Regression (minimum-norm interpolation)

$d = O(1)$ regime

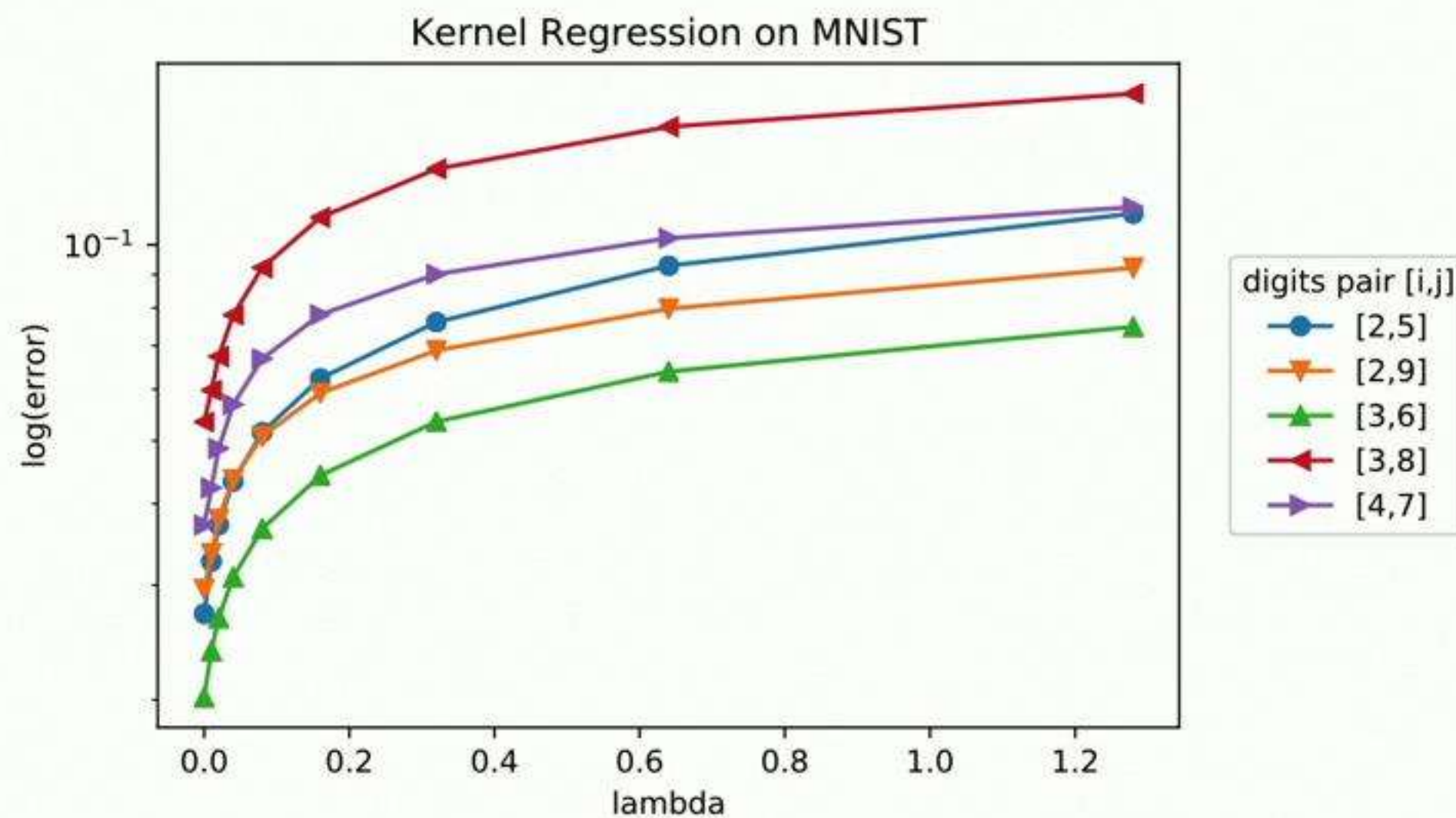
$d \asymp n^\alpha$ regime, $\alpha \in (0, 1)$

$d \asymp n$ regime

Application to Wide Neural Networks

Extra

AN EMPIRICAL EXAMPLE



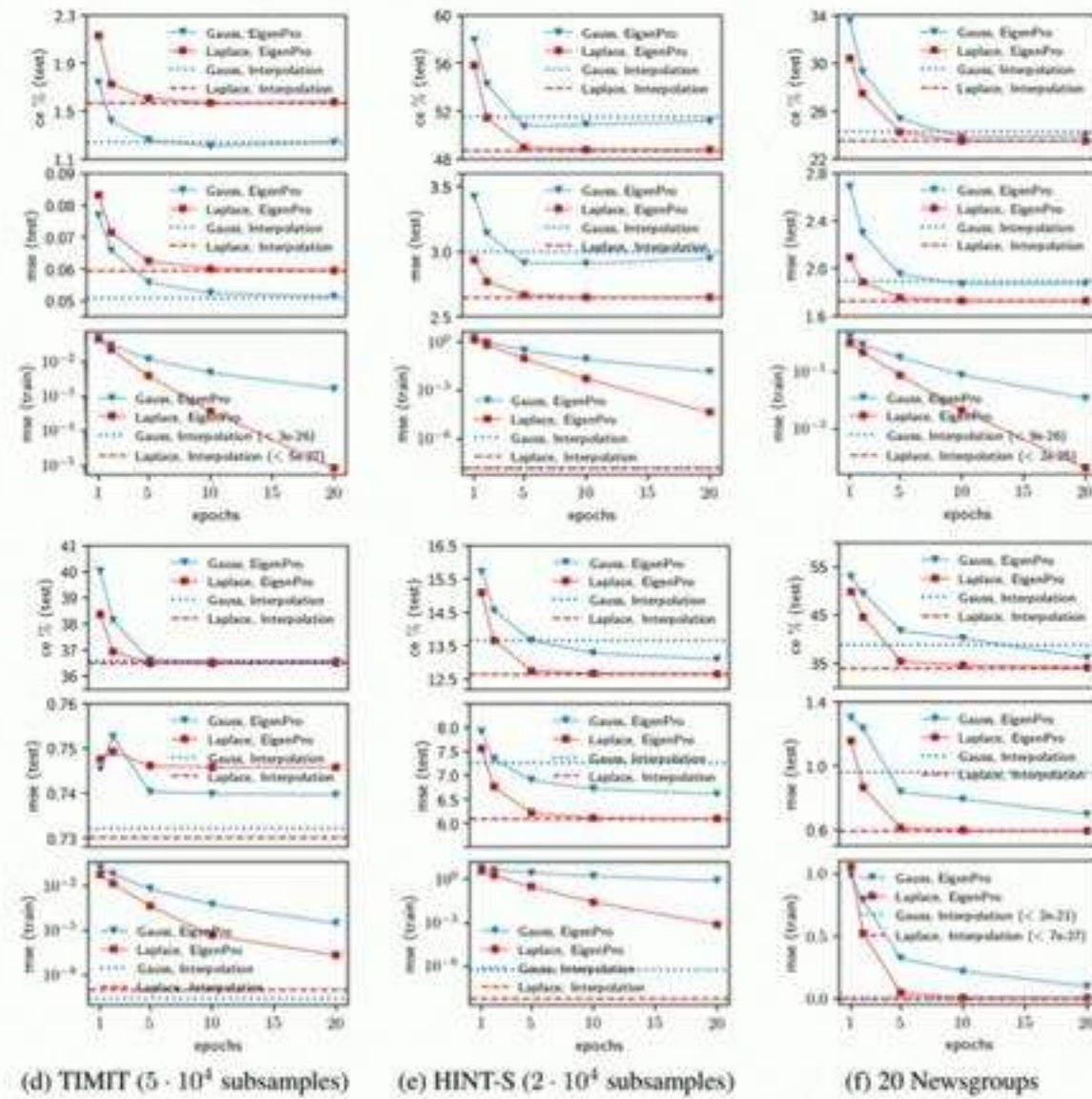
$\lambda = 0$: the interpolated solution, perfect fit on training data.

ISOLATED PHENOMENON? No

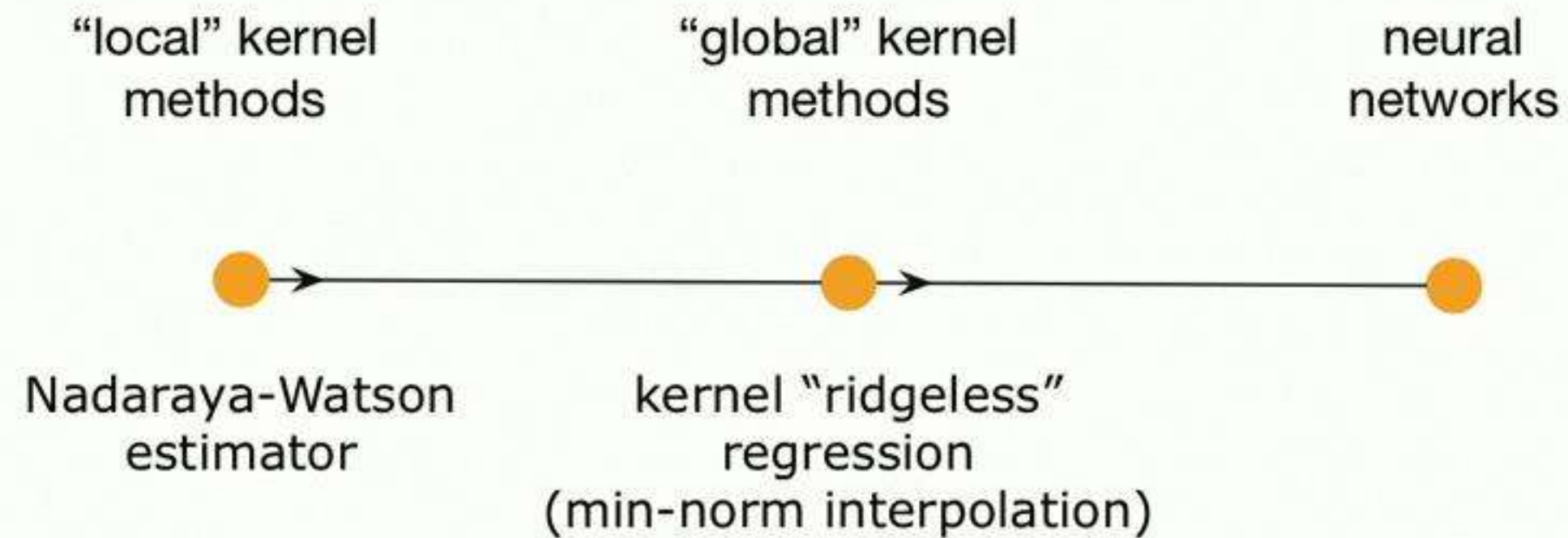
To Understand Deep Learning We Need to Understand Kernel Learning

Mikhail Belkin, Siyuan Ma, Soumik Mandal
Department of Computer Science and Engineering
Ohio State University

[mbelkin, masi]@cse.ohio-state.edu, mandal.32@osu.edu



A STUDY OF GENERALIZATION IN THE MEMORIZATION REGIME



OUTLINE

Warmup: The Nadaraya-Watson Estimator

Kernel Ridgeless Regression (minimum-norm interpolation)

$d = O(1)$ regime

$d \asymp n^\alpha$ regime, $\alpha \in (0, 1)$

$d \asymp n$ regime

Application to Wide Neural Networks

Extra

Data $Y_i = f^*(X_i) + \epsilon_i$, $i = 1, \dots, n$, and f^* is “smooth”.

Classical Nadaraya-Watson estimator:

$$\hat{f}_n(x) = \sum_{i=1}^n Y_i W_i(x)$$

where

$$W_i(x) = \frac{K((x - X_i)/h)}{\sum_{i=1}^n K((x - X_i)/h)}.$$

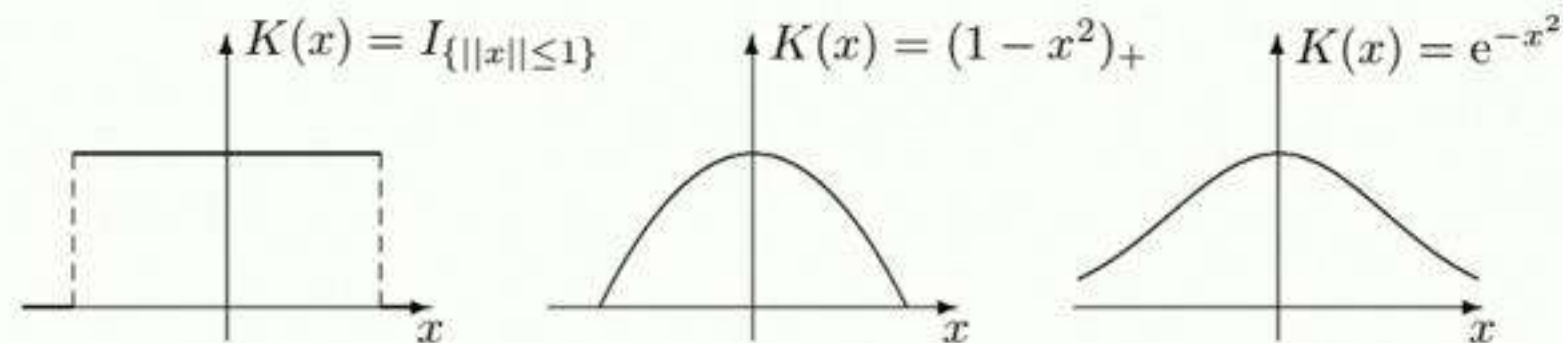


Figure 5.1. Examples for univariate kernels.

NB: “fit to data” controlled by kernel height at 0.

Example:

$$K(x) = \min\{1/\|x\|^a, \tau\}$$

We get $\widehat{f}_n(X_i) \rightarrow Y_i$ as $\tau \rightarrow \infty$.

Data $Y_i = f^*(X_i) + \epsilon_i$, $i = 1, \dots, n$, and f^* is “smooth”.

Classical Nadaraya-Watson estimator:

$$\hat{f}_n(x) = \sum_{i=1}^n Y_i W_i(x)$$

where

$$W_i(x) = \frac{K((x - X_i)/h)}{\sum_{i=1}^n K((x - X_i)/h)}.$$

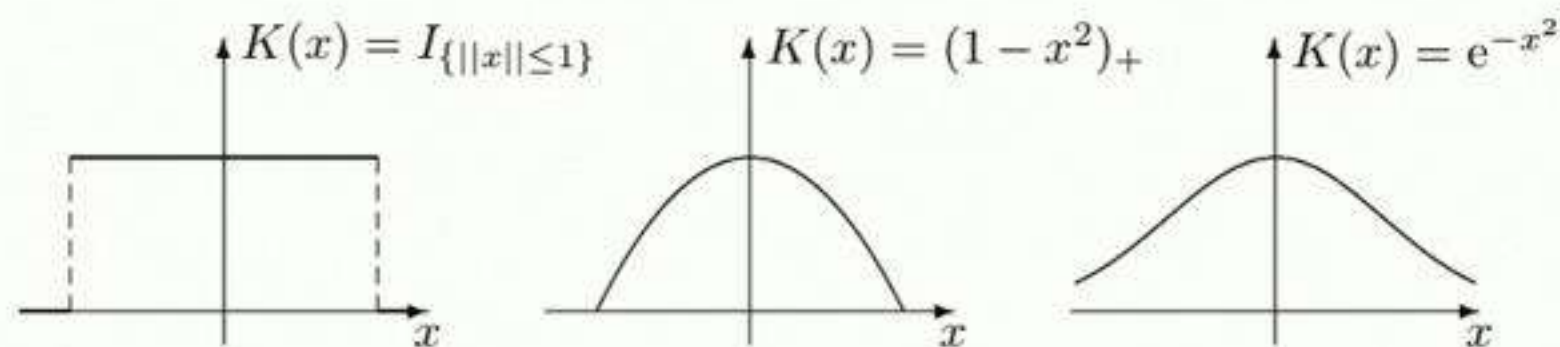


Figure 5.1. Examples for univariate kernels.

NB: "fit to data" controlled by kernel height at 0.

Example:

$$K(\mathbf{x}) = \min\{1/\|\mathbf{x}\|^d, \tau\}$$

We get $\widehat{f}_n(X_i) \rightarrow Y_i$ as $\tau \rightarrow \infty$.

{Devroye et al. (1998)}: Nadaraya-Watson with kernel $K(\mathbf{x}) = 1/\|\mathbf{x}\|^d$ is

- consistent: $|\widehat{f}_n - f^*| \rightarrow 0$ in probability (if $\mathbf{x}$ has density).
- pointwise convergence occurs nowhere: there exists a distribution such that for all $\mathbf{x}$ in support, $\widehat{f}_n(\mathbf{x})$ does not converge to $f^*(\mathbf{x})$ almost surely as $n \rightarrow \infty$.

Problem: kernel is too global (h cancels) and bias cannot be controlled.

Our fix: truncate kernel as $K(u) \triangleq \|u\|^{-\alpha} \mathbf{I}\{\|u\| \leq 1\}$.

Informal theorem: can choose bandwidth h in a minimax-optimal way such that

$$\mathbb{E} \|\widehat{f}_n - f^*\|^2 \leq C n^{-\frac{2\beta}{2\beta+d}},$$

under the assumption that f^* is β -Hölder. Note that $\widehat{f}_n$ is interpolating.

Summary:

- ▶ Fit to data does not necessarily say **anything** about overfitting.
- ▶ Data fitting part (governed by parameter τ) is decoupled from the bias-variance trade-off (as given by parameter h).

CLASSICAL PICTURE: U-SHAPED CURVE

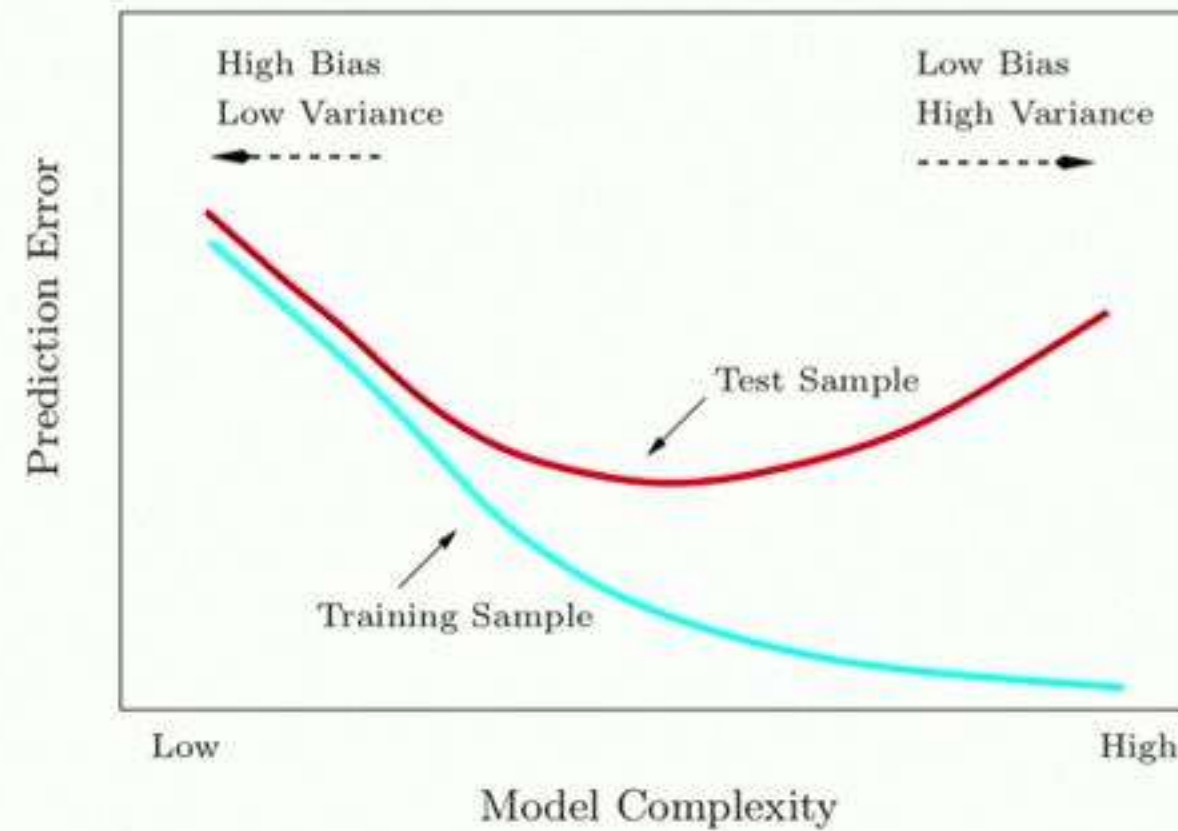


FIGURE 2.11. Test and training error as a function of model complexity.

Our fix: truncate kernel as $K(u) \triangleq \|u\|^{-\alpha} \mathbf{I}\{\|u\| \leq 1\}$.

Informal theorem: can choose bandwidth h in a minimax-optimal way such that

$$\mathbb{E} \|\widehat{f}_n - f^*\|^2 \leq C n^{-\frac{2\beta}{2\beta+d}},$$

under the assumption that f^* is β -Hölder. Note that $\widehat{f}_n$ is interpolating.

Summary:

- ▶ Fit to data does not necessarily say **anything** about overfitting.
- ▶ Data fitting part (governed by parameter τ) is decoupled from the bias-variance trade-off (as given by parameter h).

CLASSICAL PICTURE: U-SHAPED CURVE

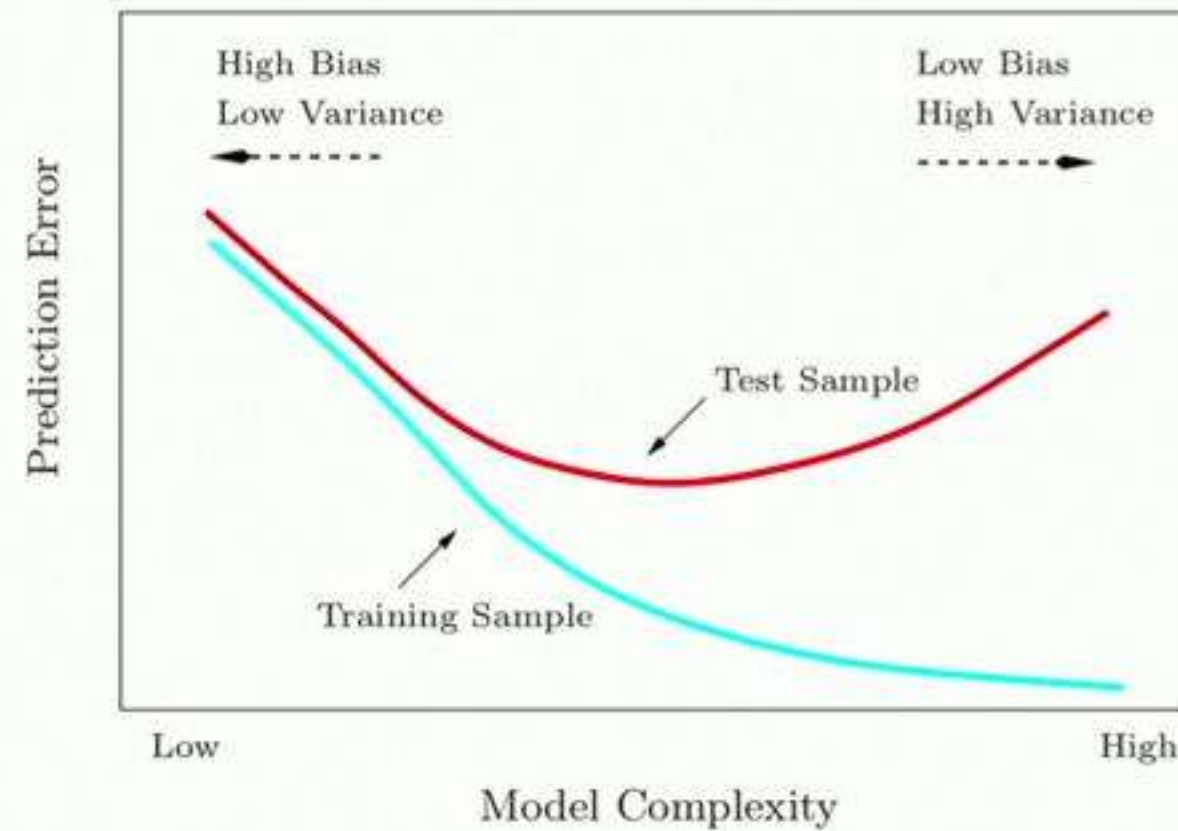
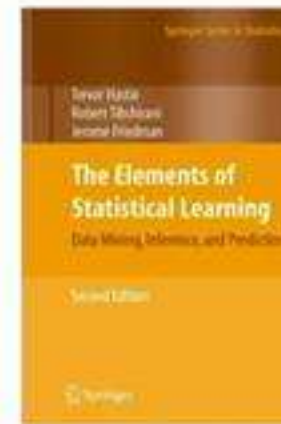


FIGURE 2.11. Test and training error as a function of model complexity.

OUTLINE

Warmup: The Nadaraya-Watson Estimator

Kernel Ridgeless Regression (minimum-norm interpolation)

$d = O(1)$ regime

$d \asymp n^\alpha$ regime, $\alpha \in (0, 1)$

$d \asymp n$ regime

Application to Wide Neural Networks

Extra

Kernel Ridge Regression:

$$\min_{f \in \mathcal{H}} \frac{1}{n} \sum_{i=1}^n (f(x_i) - y_i)^2 + \lambda \|f\|_{\mathcal{H}}^2 .$$

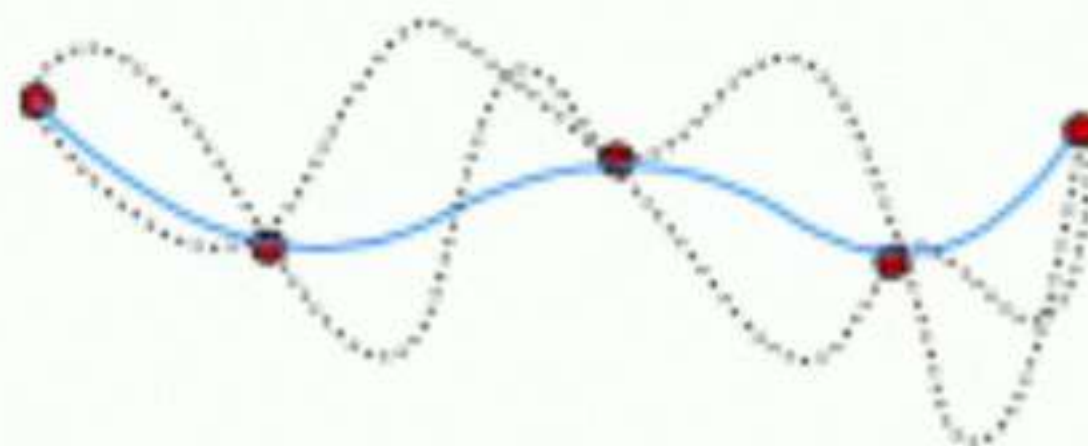
Conventional wisdom: explicit regularization $\lambda \neq 0$ added when the Reproducing Kernel Hilbert Space $\mathcal{H}$ is high- or infinite-dimensional

Puzzle: **Interpolated solutions** ($\lambda = 0$) often perform very well in practice

Belkin, Ma & Mandal (2018), Zhang, Bengio, Hardt, Recht & Vinyals (2016).

Min-norm interpolation:

$$\hat{f}_n := \operatorname{argmin}_{f \in \mathcal{H}} \|f\|_{\mathcal{H}}, \text{ s.t. } f(x_i) = y_i, \forall i \leq n.$$



- “Simplicity” is still enforced, so it is conceivable that minimum-norm interpolant generalizes.
- GD/SGD started from $\mathbf{0}$ converges to minimum-norm solution.
- Also GD for wide neural networks can be related to kernel ridgeless regression.

OUTLINE

Warmup: The Nadaraya-Watson Estimator

Kernel Ridgeless Regression (minimum-norm interpolation)

$d = O(1)$ regime

$d \asymp n^\alpha$ regime, $\alpha \in (0, 1)$

$d \asymp n$ regime

Application to Wide Neural Networks

Extra

(RAKHLIN AND ZHAI '18)

Take Laplace kernel

$$K_{\sigma}(x, x') = \sigma^{-d} \exp\{-\|x - x'\|/\sigma\}$$

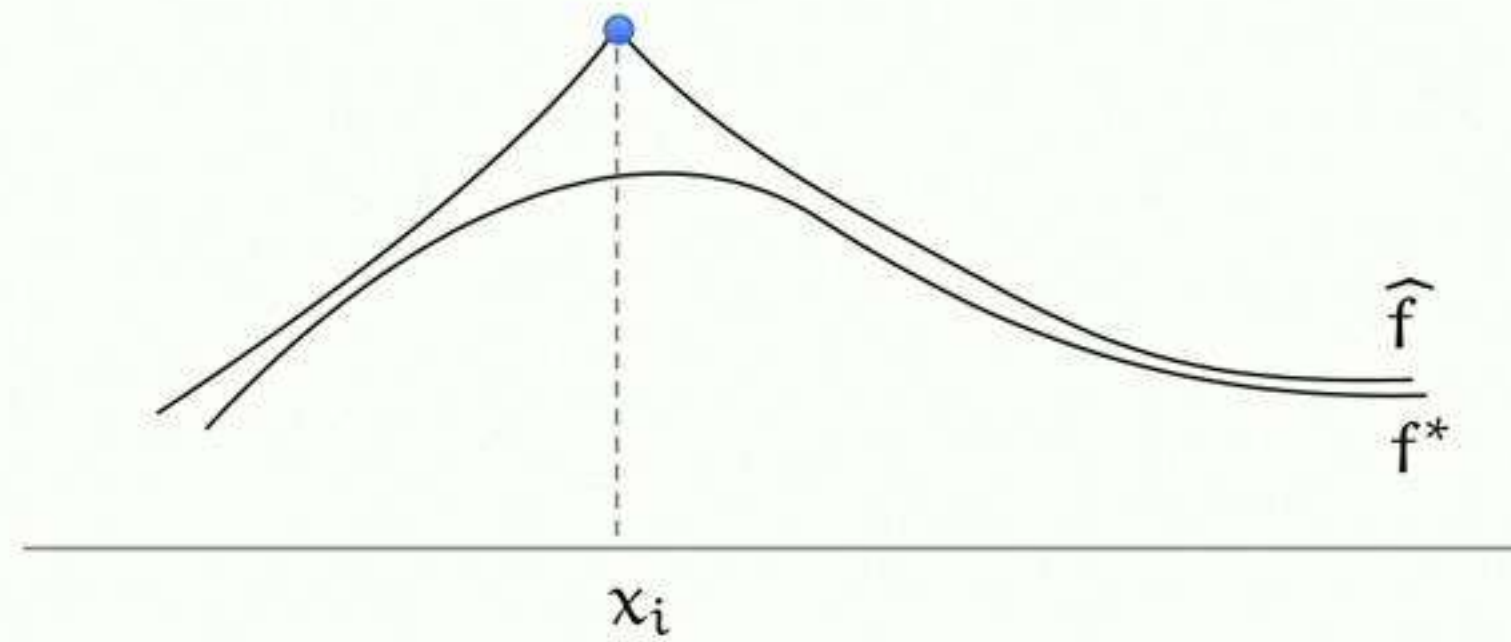
$\widehat{f}_n$ is minimum norm interpolation, as before.

Theorem: for odd dimension d , with probability $1 - O(n^{-1/2})$, for any choice of σ ,

$$\mathbb{E} \|\widehat{f}_n - f^*\|^2 \geq \Omega_d(1).$$

Hence, interpolation with Laplace kernel does not work in constant d .

LOWER BOUND FOR LARGE σ



(RAKHLIN AND ZHAI '18)

Take Laplace kernel

$$K_{\sigma}(x, x') = \sigma^{-d} \exp\{-\|x - x'\|/\sigma\}$$

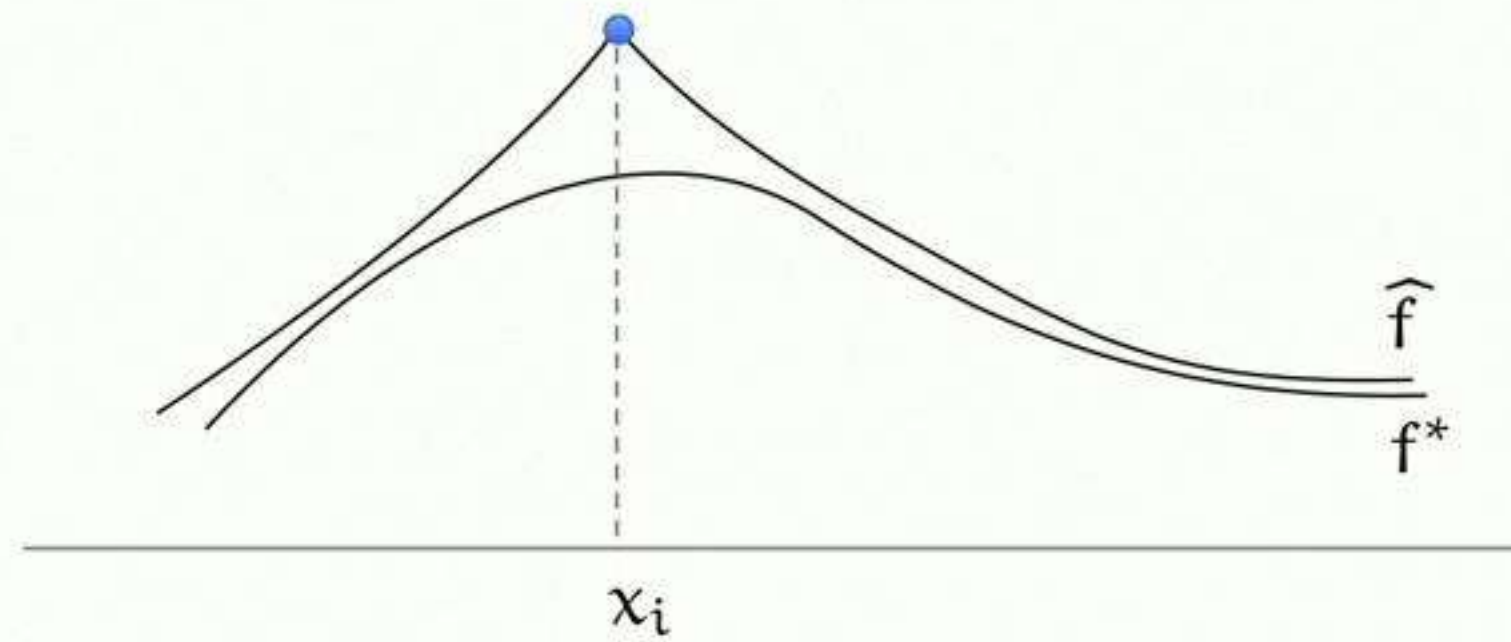
$\widehat{f}_n$ is minimum norm interpolation, as before.

Theorem: for odd dimension d , with probability $1 - O(n^{-1/2})$, for any choice of σ ,

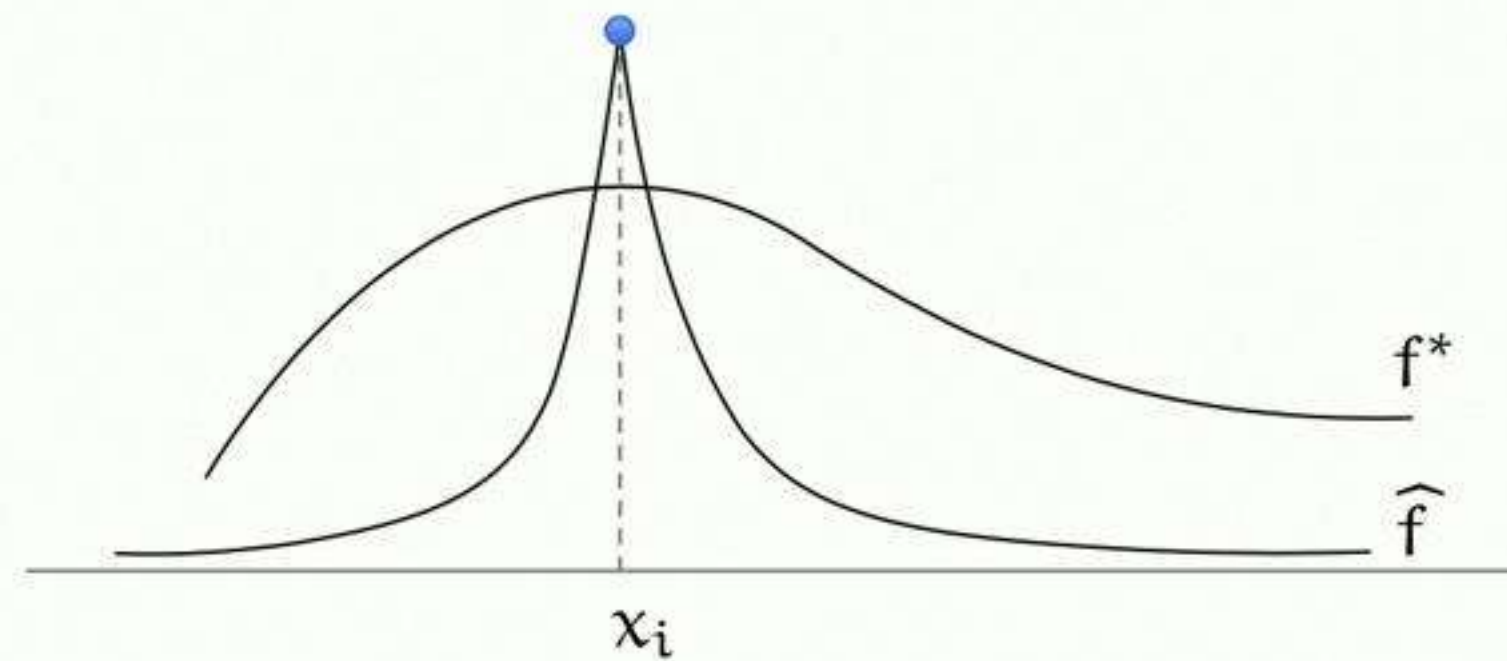
$$\mathbb{E} \|\widehat{f}_n - f^*\|^2 \geq \Omega_d(1).$$

Hence, interpolation with Laplace kernel does not work in constant d .

LOWER BOUND FOR LARGE σ



LOWER BOUND FOR SMALL σ



Interestingly, the two lower bounds cover the range of all σ , showing that Laplace interpolation cannot succeed in low dimension.

OUTLINE

Warmup: The Nadaraya-Watson Estimator

Kernel Ridgeless Regression (minimum-norm interpolation)

$d = O(1)$ regime

$d \asymp n^\alpha$ regime, $\alpha \in (0, 1)$

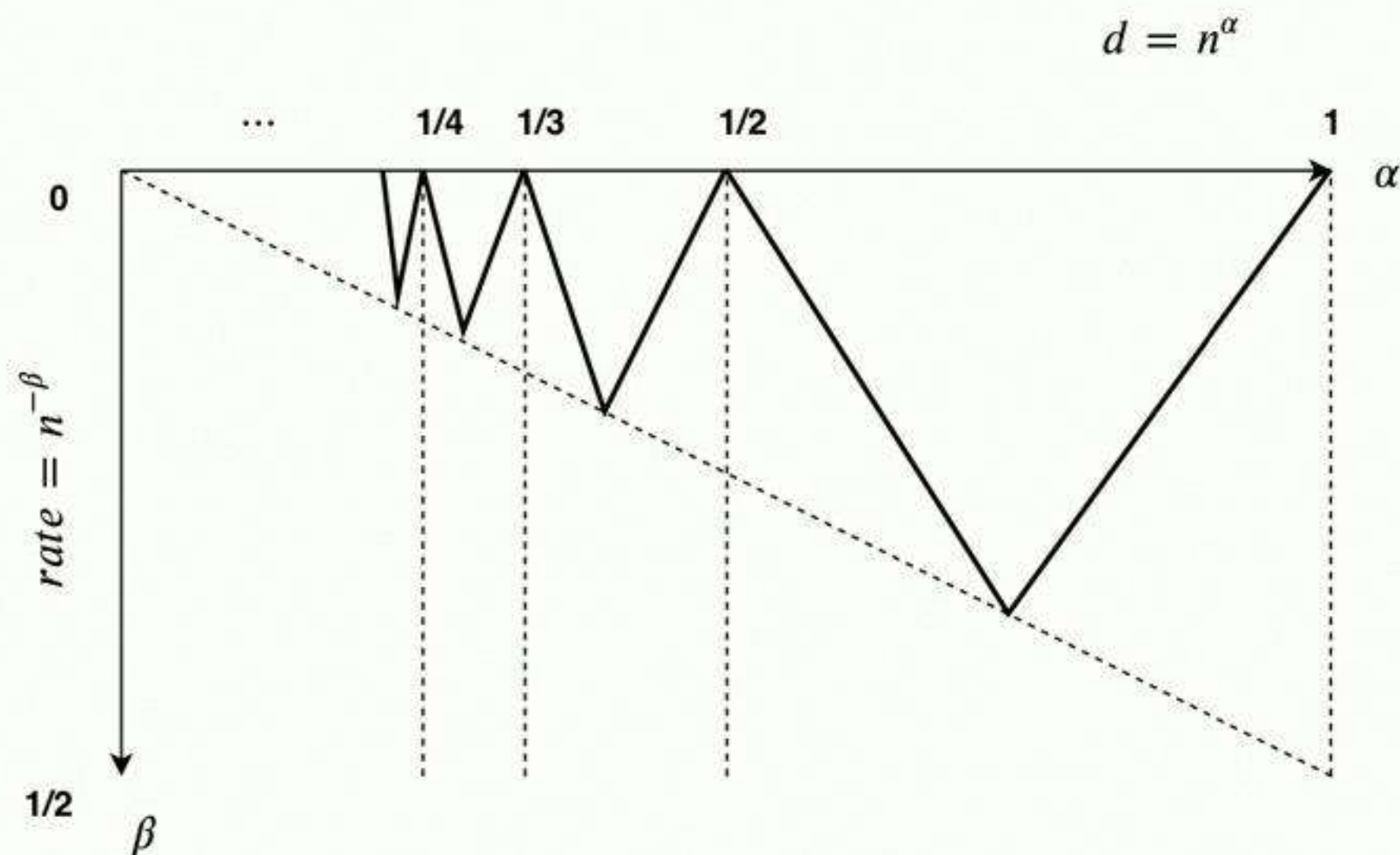
$d \asymp n$ regime

Application to Wide Neural Networks

Extra

(LIANG, RAKHLIN, ZHAI '19)

Upper bound on risk of minimum-norm interpolant has “multiple-descent” shape.



Note: we do not make assumptions on low effective rank of data (unlike (Bartlett et al '19) and (Liang and R. '18)).

Assumptions:

- ▶ Kernel is of the form $k(\mathbf{x}, \mathbf{x}') = h(\langle \mathbf{x}, \mathbf{x}' \rangle / d)$
- ▶ $\mathbf{x}_1, \dots, \mathbf{x}_n \in \mathbb{R}^d$ are i.i.d. from $\mathcal{P}^{\otimes d}$ and $\mathcal{P}$ is sub-Gaussian
- ▶ Taylor expansion

$$h(z) = \sum_{i=0}^{\infty} \alpha_i z^i \quad (1)$$

converges for all $z \in \mathbb{R}$ with all coefficients $\alpha_i > 0$.

- ▶ Regression function $f_*(\mathbf{x}) = \mathbb{E}[Y|X = \mathbf{x}]$ satisfies

$$f_*(\mathbf{x}) = \int K(\mathbf{x}, z) \rho_*(z) dz \quad (2)$$

with $\rho_*(z)$ bounded.

(LIANG, RAKHLIN, ZHAI '19)

Let $d = n^\alpha$ with $\alpha \in [\frac{1}{k_0+1}, \frac{1}{k_0}]$ for integer $k_0 \geq 1$.

Theorem (Informal).

With probability with high probability over draw of $\mathbf{X}$,

$$\mathbb{E} \left[\|\widehat{f} - f_*\|_{\mathcal{P}_X}^2 \mid \mathbf{X} \right] \leq C \cdot \left(\frac{d^{k_0}}{n} + \frac{n}{d^{k_0+1}} \right) \asymp n^{-\beta}$$

where $\beta := \min \{ (k_0 + 1)\alpha - 1, 1 - k_0\alpha \}$.

BIAS-VARIANCE

Min-norm interpolant has closed form

$$f(x) = K(x, \mathbf{X})K(\mathbf{X})^{-1}Y$$

where $K(x, \mathbf{X}) = [K(x, x_1), \dots, K(x, x_n)]$ and $[K(\mathbf{X})]_{i,j} = K(x_i, x_j)$ is the kernel matrix.

Bias-variance decomposition, conditionally on $\mathbf{X}$

$$\mathbb{E}_Y \left[\|\widehat{f}_n - f_*\|_{\mathcal{P}_X}^2 \right] = \underbrace{\mathbb{E}_Y \left[\|\widehat{f}_n - \mathbb{E}_Y[\widehat{f}_n]\|_{\mathcal{P}_X}^2 \right]}_{\text{variance}} + \underbrace{\|\mathbb{E}_Y[\widehat{f}_n] - f_*\|_{\mathcal{P}_X}^2}_{\text{bias}^2} \quad (3)$$

Up to $\text{var}(Y)$, the variance term is

$$\mathbb{E}_x \|K(x, \mathbf{X})K(\mathbf{X})^{-1}\|^2$$

and under our assumptions bias^2 is dominated by variance.

Assumptions:

- ▶ Kernel is of the form $k(\mathbf{x}, \mathbf{x}') = h(\langle \mathbf{x}, \mathbf{x}' \rangle / d)$
- ▶ $\mathbf{x}_1, \dots, \mathbf{x}_n \in \mathbb{R}^d$ are i.i.d. from $\mathcal{P}^{\otimes d}$ and $\mathcal{P}$ is sub-Gaussian
- ▶ Taylor expansion

$$h(z) = \sum_{i=0}^{\infty} \alpha_i z^i \quad (1)$$

converges for all $z \in \mathbb{R}$ with all coefficients $\alpha_i > 0$.

- ▶ Regression function $f_*(\mathbf{x}) = \mathbb{E}[Y|X = \mathbf{x}]$ satisfies

$$f_*(\mathbf{x}) = \int K(\mathbf{x}, z) \rho_*(z) dz \quad (2)$$

with $\rho_*(z)$ bounded.

BIAS-VARIANCE

Min-norm interpolant has closed form

$$f(x) = K(x, \mathbf{X})K(\mathbf{X})^{-1}Y$$

where $K(x, \mathbf{X}) = [K(x, x_1), \dots, K(x, x_n)]$ and $[K(\mathbf{X})]_{i,j} = K(x_i, x_j)$ is the kernel matrix.

Bias-variance decomposition, conditionally on $\mathbf{X}$

$$\mathbb{E}_Y \left[\|\widehat{f}_n - f_*\|_{\mathcal{P}_X}^2 \right] = \underbrace{\mathbb{E}_Y \left[\|\widehat{f}_n - \mathbb{E}_Y[\widehat{f}_n]\|_{\mathcal{P}_X}^2 \right]}_{\text{variance}} + \underbrace{\|\mathbb{E}_Y[\widehat{f}_n] - f_*\|_{\mathcal{P}_X}^2}_{\text{bias}^2} \quad (3)$$

Up to $\text{var}(Y)$, the variance term is

$$\mathbb{E}_x \|K(x, \mathbf{X})K(\mathbf{X})^{-1}\|^2$$

and under our assumptions bias^2 is dominated by variance.

ROADMAP OF THE PROOF (I):

Consider truncated kernel

$$h^{[k]}(x) = \sum_{i=0}^k \alpha_i x^i, \quad K_{ij}^{[k]} := h^{[k]}(\langle x_i, x_j \rangle / d) / n.$$

ROADMAP OF THE PROOF (I):

Consider truncated kernel

$$h^{[k]}(x) = \sum_{i=0}^k \alpha_i x^i, \quad K_{ij}^{[k]} := h^{[k]}(\langle x_i, x_j \rangle / d) / n.$$

Proposition (Restricted Lower Isometry Property for Kernel).

With high prob, for any $k \leq k_0$, $K^{[k]}$ has $\binom{d+k}{k}$ nonzero eigenvalues, all of them larger than Cd^{-k} , and the range of $K^{[k]}$ is

$$\{(p(x_1), \dots, p(x_n)) : p \text{ is a multivar. polynomial of deg. } \leq k\}.$$

ROADMAP OF THE PROOF (II):

- ▶ Let $(r_1 \cdots r_d)_{k_0}$ be index in $\{1, 2, \dots, \binom{k_0+d}{d}\}$
- ▶ Taylor expansion gives an $n \times \binom{k_0+d}{d}$ matrix Φ as

$$\Phi_{i, (r_1 \cdots r_d)_{k_0}} = \sqrt{c_{r_1 \cdots r_d} \alpha_{r_1 + \cdots + r_d}} \mathbf{p}_{r_1 \cdots r_d}(x_i) / d^{(r_1 + \cdots + r_d)/2}$$

such that

$$\mathbf{K}^{[k_0]} = \frac{1}{n} \Phi \Phi^\top .$$

- ▶ Hard to analyze because of correlations.
- ▶ Gram-Schmidt Process to orthogonalize $1, z, z^2, \dots$ with respect to $L_2(\mathcal{P})$. This gives $q_1, q_2, \dots$ on $\mathbb{R}$.

$$\Psi_{i, (r_1 \cdots r_d)_{k_0}} = \sqrt{c_{r_1 \cdots r_d} \alpha_{r_1 + \cdots + r_d}} \mathbf{q}_{r_1 \cdots r_d}(x_i) / d^{(r_1 + \cdots + r_d)/2} .$$

where $\mathbf{q}_{r_1 \cdots r_d}(x) := \prod_{i=1}^d q_{r_i}(x[i])$.

- ▶ Prove that this process gives $\Phi = \Psi \Lambda$ with bounded Λ in operator norm.

ROADMAP OF THE PROOF (III):

- ▶ Let $f_\gamma(x) := \sum_{\leq k_0} \gamma_{r_1 \dots r_d} q_{r_1 \dots r_d}(x)$, then

$$\mathbb{E} f_\gamma(x)^4 \lesssim (\mathbb{E} f_\gamma(x)^2)^2$$

- ▶ Show a “small-ball” property for functions

$$f_u(x) = \sum_{r_1, \dots, r_d \geq 0, \sum_i r_i \leq k_0} u_{(r_1 \dots r_d)_{k_0}} \sqrt{c_{r_1 \dots r_d} \alpha_{r_1 + \dots + r_d}} q_{r_1 \dots r_d}(x) / d^{(r_1 + \dots + r_d)/2} .$$

That is, there exist θ, δ , such that for any u

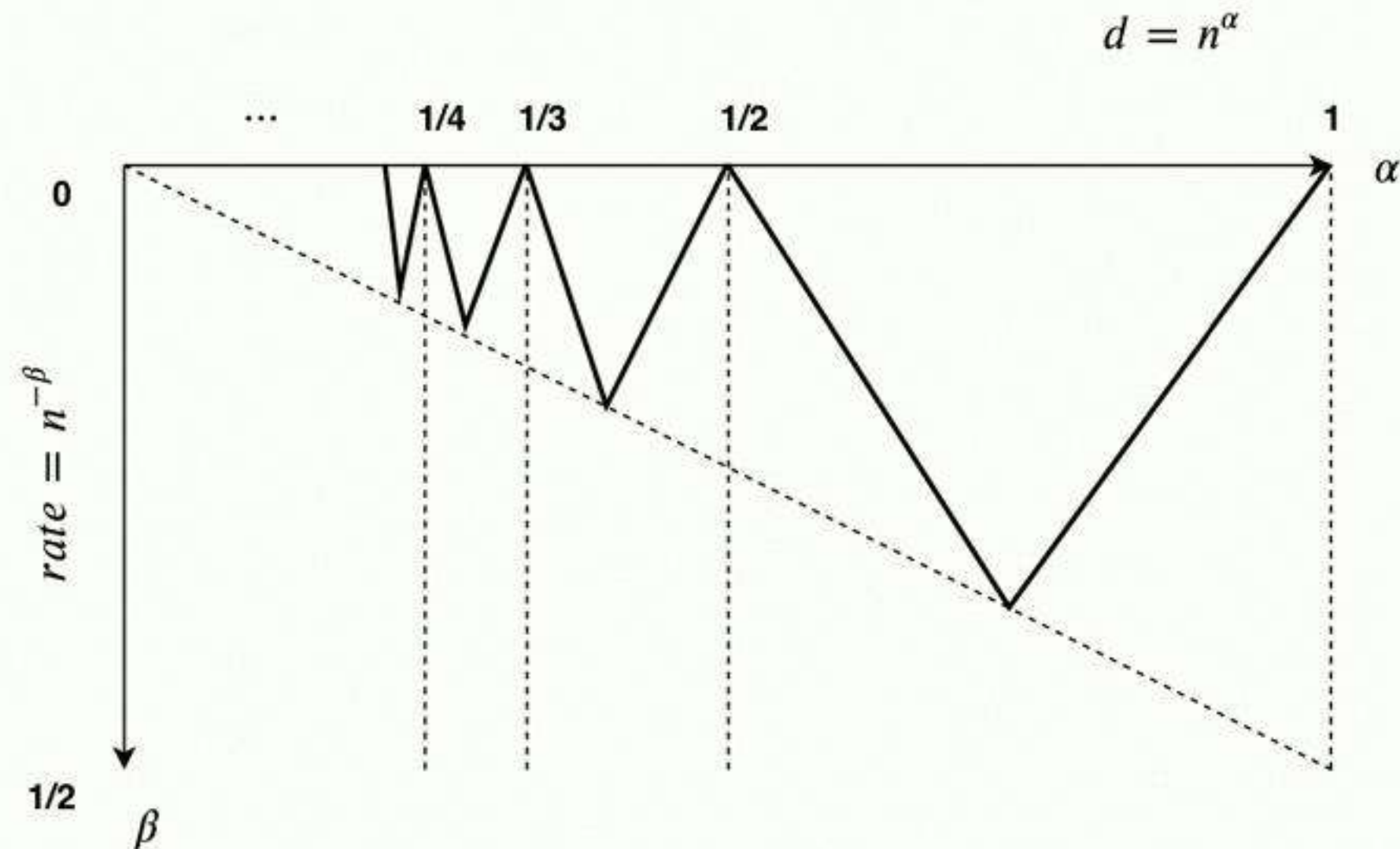
$$\mathbb{P}(f_u(x)^2 > \theta \mathbb{E}[f_u(x)^2]) \geq \delta .$$

- ▶ Lower bound eigen-values of

$$\Psi^\top \Psi / n > C d^{-k_0}$$

(LIANG, RAKHLIN, ZHAI '19)

Upper bound on risk of minimum-norm interpolant has “multiple-descent” shape.



OUTLINE

Warmup: The Nadaraya-Watson Estimator

Kernel Ridgeless Regression (minimum-norm interpolation)

$d = O(1)$ regime

$d \asymp n^\alpha$ regime, $\alpha \in (0, 1)$

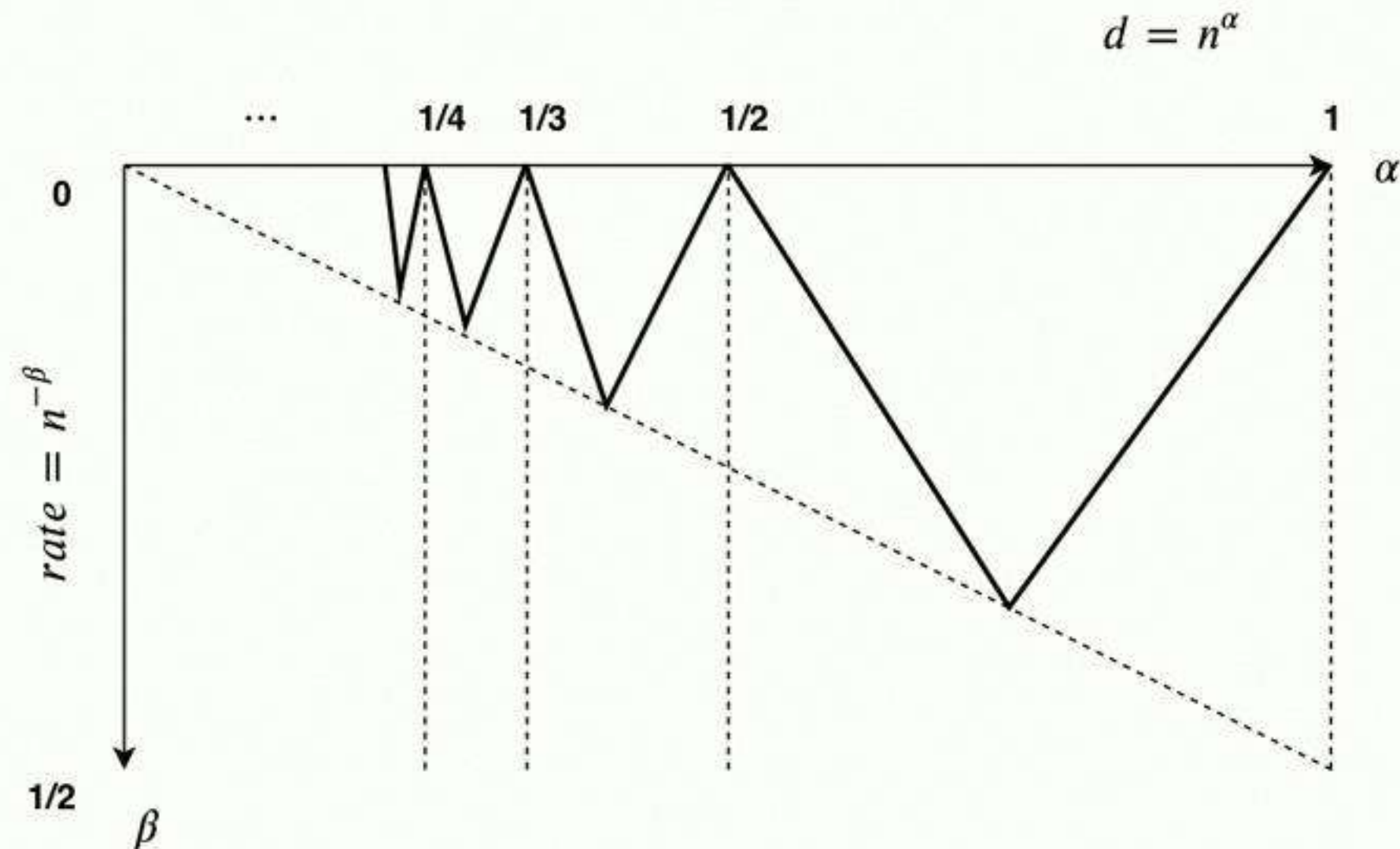
$d \asymp n$ regime

Application to Wide Neural Networks

Extra

(LIANG, RAKHLIN, ZHAI '19)

Upper bound on risk of minimum-norm interpolant has “multiple-descent” shape.



ASSUMPTIONS

- ▶ (A.1) High-dim regime: both d, n large, $c \leq d/n \leq C$
 $\Sigma_d = \mathbf{E}_\mu[\mathbf{x}\mathbf{x}^*]$ satisfies $\|\Sigma_d\| \leq C$ and $\text{tr}(\Sigma_d)/d \geq c$.
- ▶ (A.2) $(8 + m)$ -moments: $\mathbf{z}_i := \Sigma_d^{-1/2} \mathbf{x}_i$
each entries of $\mathbf{z}_i$ are i.i.d. mean zero, with bounded $(8 + m)$ -moments.
- ▶ (A.3) Noise condition: $\mathbf{E}[(f_*(\mathbf{x}) - \mathbf{y})^2 | \mathbf{x} = \mathbf{x}] \leq \sigma^2$ for all $\mathbf{x} \in \Omega$.
- ▶ (A.4) Non-linear kernel: for a non-linear smooth function $h(\cdot)$

$$K(\mathbf{x}, \mathbf{x}') = h\left(\frac{1}{d} \langle \mathbf{x}, \mathbf{x}' \rangle\right)$$

NB: (A.4) can be replaced with $h\left(\frac{1}{d} \|\mathbf{x} - \mathbf{x}'\|^2\right)$.

OUTLINE

Warmup: The Nadaraya-Watson Estimator

Kernel Ridgeless Regression (minimum-norm interpolation)

- $d = O(1)$ regime

- $d \asymp n^\alpha$ regime, $\alpha \in (0, 1)$

- $d \asymp n$ regime

Application to Wide Neural Networks

Extra

One-hidden-layer NN

$$f(\mathbf{x}; W, \mathbf{a}) := \frac{1}{\sqrt{m}} \sum_{i=1}^m a_i \sigma(\mathbf{w}_i^\top \tilde{\mathbf{x}}), \quad (5)$$

where $W = (\mathbf{w}_1, \dots, \mathbf{w}_m) \in \mathbb{R}^{(d+1) \times m}$ matrix and $\mathbf{a} \in \mathbb{R}^m$ and $\tilde{\mathbf{x}} = (\mathbf{x}^\top, \sqrt{d})^\top$.

Square loss:

$$L = \frac{1}{2n} \sum_{j=1}^n (f(\mathbf{x}_j; W, \mathbf{a}) - y_j)^2 \quad (6)$$

Gradient:

$$\frac{\partial L}{\partial a_i} = \frac{1}{n} \sum_{j=1}^n \frac{\sigma(\mathbf{w}_i^\top \tilde{\mathbf{x}}_j)}{\sqrt{m}} (f(\mathbf{x}_j; W, \mathbf{a}) - y_j), \quad (7)$$

and

$$\frac{\partial L}{\partial \mathbf{w}_i} = \frac{1}{n} \sum_{j=1}^n \frac{a_i \tilde{\mathbf{x}}_j \sigma'(\mathbf{w}_i^\top \tilde{\mathbf{x}}_j)}{\sqrt{m}} (f(\mathbf{x}_j; W, \mathbf{a}) - y_j). \quad (8)$$

Gradient flow yields

$$\frac{df(\mathbf{x}; W(t), \mathbf{a}(t))}{dt} = -\frac{1}{n} \sum_{j=1}^n h^m(\mathbf{x}, \mathbf{x}_j) (f(\mathbf{x}_j; W, \mathbf{a}) - y_j)$$

for

$$h^m(\mathbf{x}, \mathbf{x}') := \frac{1}{m} \left(\sum_{i=1}^m \sigma(\mathbf{w}_i^\top \tilde{\mathbf{x}}) \sigma(\mathbf{w}_i^\top \tilde{\mathbf{x}}') + \mathbf{x}^\top \mathbf{x}' \sum_{i=1}^m a_i^2 \sigma'(\mathbf{w}_i^\top \tilde{\mathbf{x}}) \sigma'(\mathbf{w}_i^\top \tilde{\mathbf{x}}') \right) \quad (9)$$

and $\mathbf{f}$ converges to the minimum-norm interpolant.

With $N(0, 1)$ independent init, as $m \rightarrow \infty$, H^m converges to

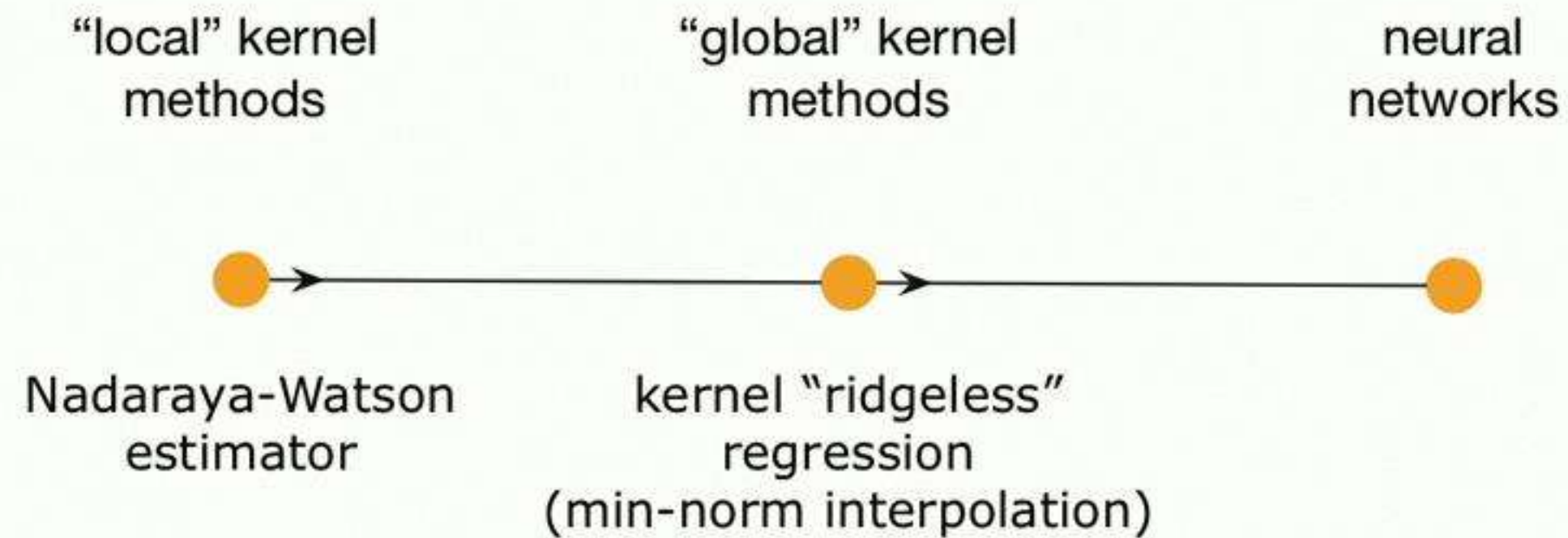
$$h^\infty(x, x') := \frac{1}{4\pi} \|\tilde{x}\| \|\tilde{x}'\| \mathcal{U}(\cos \theta_{\tilde{x}, \tilde{x}'})$$

where $\theta_{\tilde{x}, \tilde{x}'}$ is the angle between $\tilde{x}$ and $\tilde{x}'$ and

$$\mathcal{U}(t) := 3t(\pi - \arccos(t)) + \sqrt{1 - t^2}$$

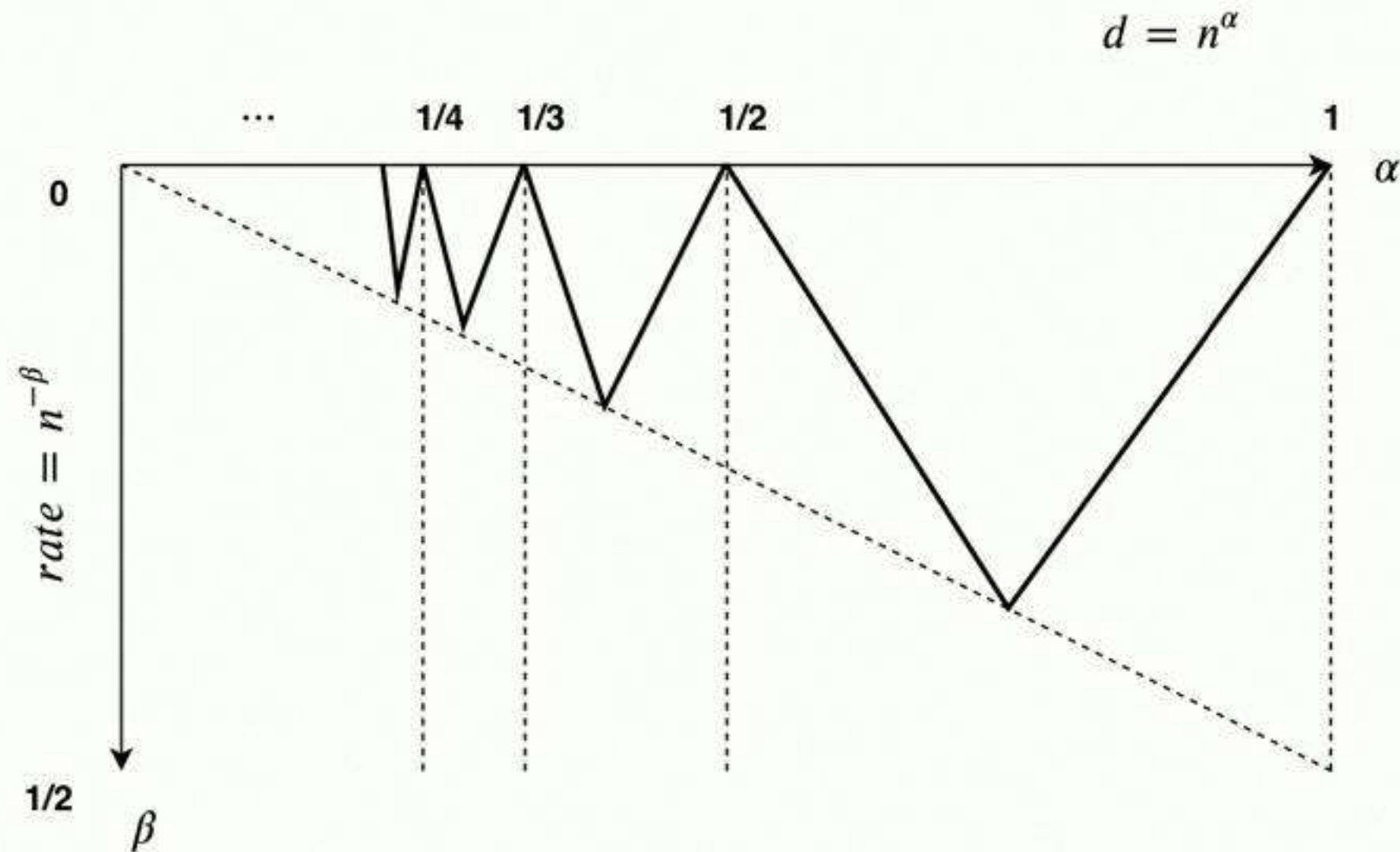
The theorem stated earlier in the talk for the $d \asymp n^\alpha$ regime can be extended to this kernel.

A STUDY OF GENERALIZATION IN THE MEMORIZATION REGIME



(LIANG, RAKHLIN, ZHAI '19)

Upper bound on risk of minimum-norm interpolant has “multiple-descent” shape.



Assumptions:

- ▶ Kernel is of the form $k(\mathbf{x}, \mathbf{x}') = h(\langle \mathbf{x}, \mathbf{x}' \rangle / d)$
- ▶ $\mathbf{x}_1, \dots, \mathbf{x}_n \in \mathbb{R}^d$ are i.i.d. from $\mathcal{P}^{\otimes d}$ and $\mathcal{P}$ is sub-Gaussian
- ▶ Taylor expansion

$$h(z) = \sum_{i=0}^{\infty} \alpha_i z^i \quad (1)$$

converges for all $z \in \mathbb{R}$ with all coefficients $\alpha_i > 0$.

- ▶ Regression function $f_*(\mathbf{x}) = \mathbb{E}[Y|X = \mathbf{x}]$ satisfies

$$f_*(\mathbf{x}) = \int K(\mathbf{x}, z) \rho_*(z) dz \quad (2)$$

with $\rho_*(z)$ bounded.

(LIANG, RAKHLIN, ZHAI '19)

Let $d = n^\alpha$ with $\alpha \in [\frac{1}{k_0+1}, \frac{1}{k_0}]$ for integer $k_0 \geq 1$.

Theorem (Informal).

With probability with high probability over draw of $\mathbf{X}$,

$$\mathbb{E} [\|\widehat{f} - f_*\|_{\mathcal{P}_X}^2 \mid \mathbf{X}] \leq C \cdot \left(\frac{d^{k_0}}{n} + \frac{n}{d^{k_0+1}} \right) \asymp n^{-\beta}$$

where $\beta := \min \{ (k_0 + 1)\alpha - 1, 1 - k_0\alpha \}$.

Assumptions:

- ▶ Kernel is of the form $k(\mathbf{x}, \mathbf{x}') = h(\langle \mathbf{x}, \mathbf{x}' \rangle / d)$
- ▶ $\mathbf{x}_1, \dots, \mathbf{x}_n \in \mathbb{R}^d$ are i.i.d. from $\mathcal{P}^{\otimes d}$ and $\mathcal{P}$ is sub-Gaussian
- ▶ Taylor expansion

$$h(z) = \sum_{i=0}^{\infty} \alpha_i z^i \quad (1)$$

converges for all $z \in \mathbb{R}$ with all coefficients $\alpha_i > 0$.

- ▶ Regression function $f_*(\mathbf{x}) = \mathbb{E}[Y|X = \mathbf{x}]$ satisfies

$$f_*(\mathbf{x}) = \int K(\mathbf{x}, z) \rho_*(z) dz \quad (2)$$

with $\rho_*(z)$ bounded.



The generalization error of random feature models

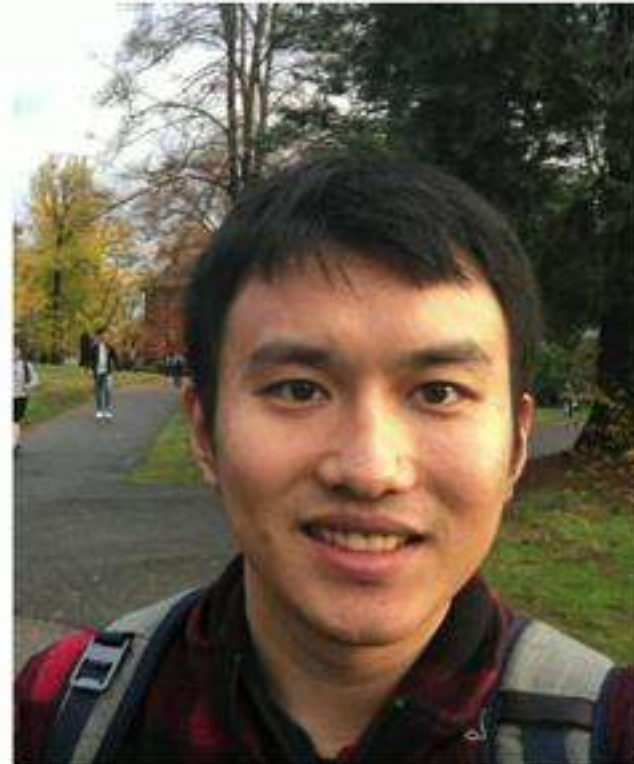
Andrea Montanari

Stanford University

August 26, 2019

3 papers, 6 co-authors

“...PRECISE ASYMPTOTICS AND DOUBLE DESCENT”



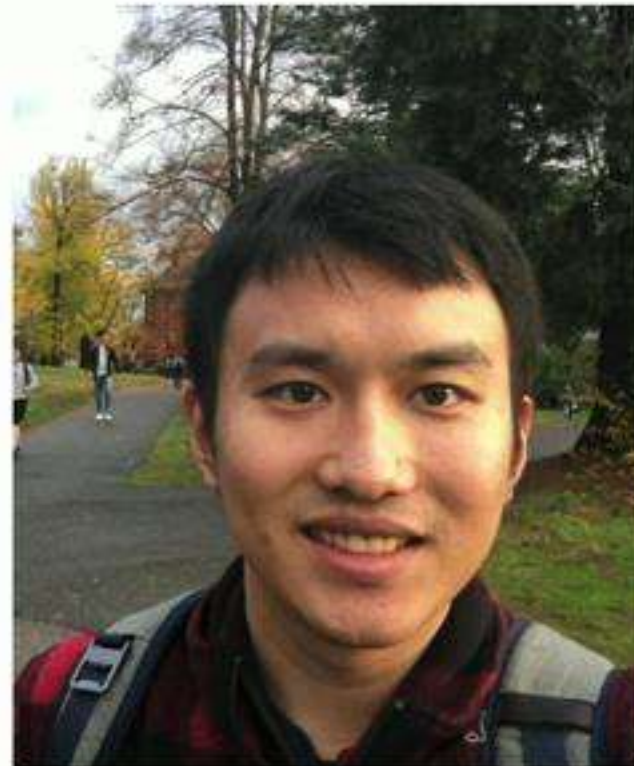
Song Mei

arXiv:1908.05355

“LINEARIZED TWO-LAYERS NEURAL NETWORKS...”



Behrooz Ghorbani



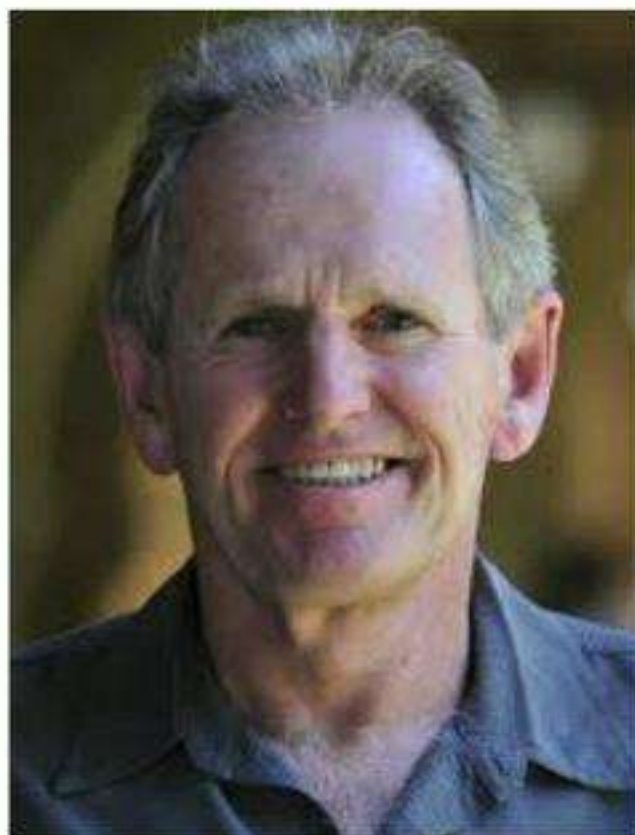
Song Mei



Theodor Misiakiewicz

arXiv:1904.12191

“SURPRISES IN HIGH-DIMENSIONAL RIDGELESS ...”



Trevor Hastie



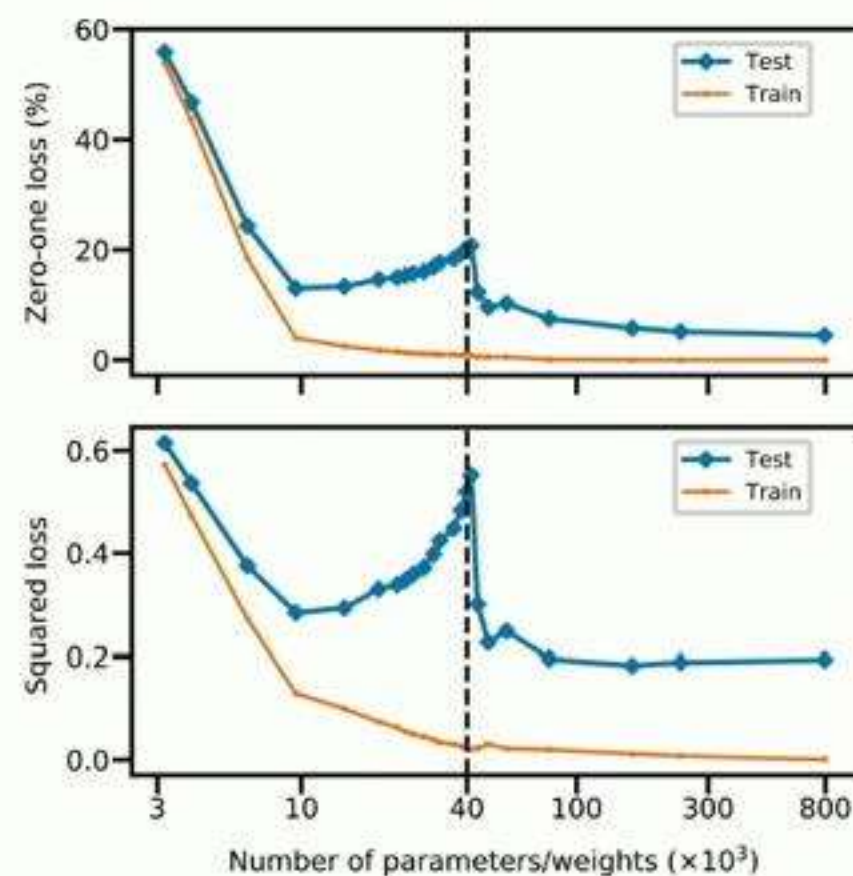
Saharon Rosset



Ryan Tibshirani

arXiv:1903.08560

Motivation



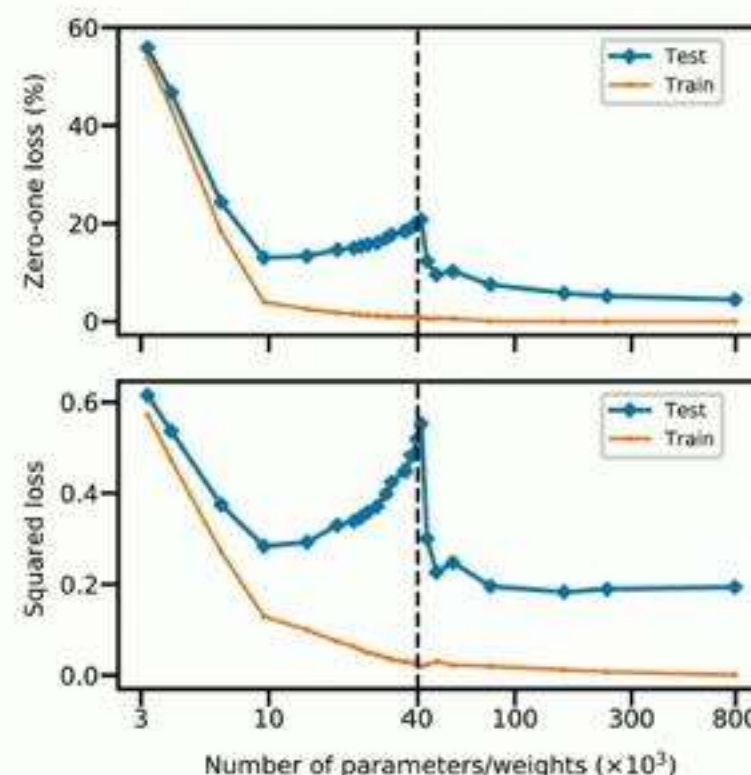
Belkin, Hsu, Mandal, 2019

Two-layers fully connected, MNIST

Do we understand these curves?

Defining the question

Desiderata



- ✓ Peak at the interpolation threshold
- ✓ Global minimum in the overparametrized regime
- ✓ Monotone decreasing in the overparametrized regime
- ✓ Vanishing (explicit) regularization

Ideally: Two-layers neural networks

$$\mathcal{F}_{\text{NN}} \equiv \left\{ f(x) = \sum_{i=1}^N a_i \sigma(\langle w_i, x \rangle) : a_i \in \mathbb{R}, w_i \in \mathbb{R}^d \forall i \leq N \right\}.$$

- ▶ Train by SGD until TrainingError= 0
- ▶ Compute TestError for this network
- ▶ Perhaps in a couple of years

Two linearizations...

$$\mathcal{F}_{\text{RF}}(\mathbf{W}) \equiv \left\{ f(\mathbf{x}) = \sum_{i=1}^N a_i \sigma(\langle \mathbf{w}_i, \mathbf{x} \rangle) : a_i \in \mathbb{R} \forall i \leq N \right\}.$$

- $\mathbf{W} = [\mathbf{w}_1, \dots, \mathbf{w}_N]$ random; $\mathbf{w}_i \sim \text{Unif}(\mathbb{S}^{d-1}(1))$, i.i.d.

$$\mathcal{F}_{\text{NT}}(\mathbf{W}) \equiv \left\{ f(\mathbf{x}) = \sum_{i=1}^N \langle \mathbf{a}_i, \mathbf{x} \rangle \sigma'(\langle \mathbf{w}_i, \mathbf{x} \rangle) : \mathbf{a}_i \in \mathbb{R}^d \forall i \leq N \right\}.$$

- ▶ $\mathbf{W} = [\mathbf{w}_1, \dots, \mathbf{w}_N]$ random; $\mathbf{w}_i \sim \text{Unif}(\mathbb{S}^{d-1}(1))$, i.i.d.
- ▶ Linear expansion of NN around random initialization

The rest of this talk

1. A stylized (Gaussian) model
2. Nonlinear features: Approximation in high dimension
3. Nonlinear features: Generalization curve

A stylized (Gaussian) model

Setting

- ▶ Data $\{(y_i, x_i)\}_{i \leq n}$, $z_i = \phi(x_i) \in \mathbb{R}^p$

$$y_i = \langle \beta_0, z_i \rangle + \varepsilon_i, \quad \varepsilon_i \sim N(0, \tau^2)$$

- ▶ ‘Ridgeless regression’

$$\hat{\beta}(\lambda) \equiv \arg \min_{\beta \in \mathbb{R}^p} \left\{ \|y - Z\beta\|_2^2 + \lambda \|\beta\|_2^2 \right\},$$

$$\hat{\beta}(0) \equiv \lim_{\lambda \rightarrow 0+} \hat{\beta}(\lambda)$$

- ▶ Test error

$$R(\beta_0) = \mathbb{E}\{(\langle \hat{\beta}, z^{\text{new}} \rangle - \langle \beta_0, z^{\text{new}} \rangle)^2\}$$

A stylized model

- ▶ Gaussian features

$$z_i \sim \mathcal{N}(\mathbf{0}, \Sigma)$$

Setting

- ▶ Data $\{(y_i, x_i)\}_{i \leq n}$, $z_i = \phi(x_i) \in \mathbb{R}^p$

$$y_i = \langle \beta_0, z_i \rangle + \varepsilon_i, \quad \varepsilon_i \sim N(0, \tau^2)$$

- ▶ ‘Ridgeless regression’

$$\hat{\beta}(\lambda) \equiv \arg \min_{\beta \in \mathbb{R}^p} \left\{ \|y - Z\beta\|_2^2 + \lambda \|\beta\|_2^2 \right\},$$

$$\hat{\beta}(0) \equiv \lim_{\lambda \rightarrow 0+} \hat{\beta}(\lambda)$$

- ▶ Test error

$$R(\beta_0) = \mathbb{E}\{(\langle \hat{\beta}, z^{\text{new}} \rangle - \langle \beta_0, z^{\text{new}} \rangle)^2\}$$

Precise asymptotics

Theorem (Hastie, M., Rosset, Tibshirani, 2019)

If $p/n \rightarrow \gamma$, $\Sigma = I_p$, $\|\beta_0\|_2 \rightarrow r^2$, then

$$R(\beta_0) \rightarrow \begin{cases} \tau^2 \frac{\gamma}{1-\gamma} & \text{for } \gamma < 1, \\ \tau^2 \frac{1}{\gamma-1} + r^2 \left(1 - \frac{1}{\gamma}\right) & \text{for } \gamma > 1. \end{cases}$$

Further:

- ▶ If $\Sigma \neq I_p$. . .
- ▶ If $z_i = \Sigma^{1/2} x_i$, x_i with iid components . . .

[Belkin, Hsu, Xu, 2019; Bartlett, Long, Lugosi, Tsigler, 2019]

Setting

- ▶ Data $\{(y_i, x_i)\}_{i \leq n}$, $z_i = \phi(x_i) \in \mathbb{R}^p$

$$y_i = \langle \beta_0, z_i \rangle + \varepsilon_i, \quad \varepsilon_i \sim N(0, \tau^2)$$

- ▶ ‘Ridgeless regression’

$$\hat{\beta}(\lambda) \equiv \arg \min_{\beta \in \mathbb{R}^p} \left\{ \|y - Z\beta\|_2^2 + \lambda \|\beta\|_2^2 \right\},$$

$$\hat{\beta}(0) \equiv \lim_{\lambda \rightarrow 0+} \hat{\beta}(\lambda)$$

- ▶ Test error

$$R(\beta_0) = \mathbb{E}\{(\langle \hat{\beta}, z^{\text{new}} \rangle - \langle \beta_0, z^{\text{new}} \rangle)^2\}$$

Precise asymptotics

Theorem (Hastie, M., Rosset, Tibshirani, 2019)

If $p/n \rightarrow \gamma$, $\Sigma = \mathbf{I}_p$, $\|\beta_0\|_2 \rightarrow r^2$, then

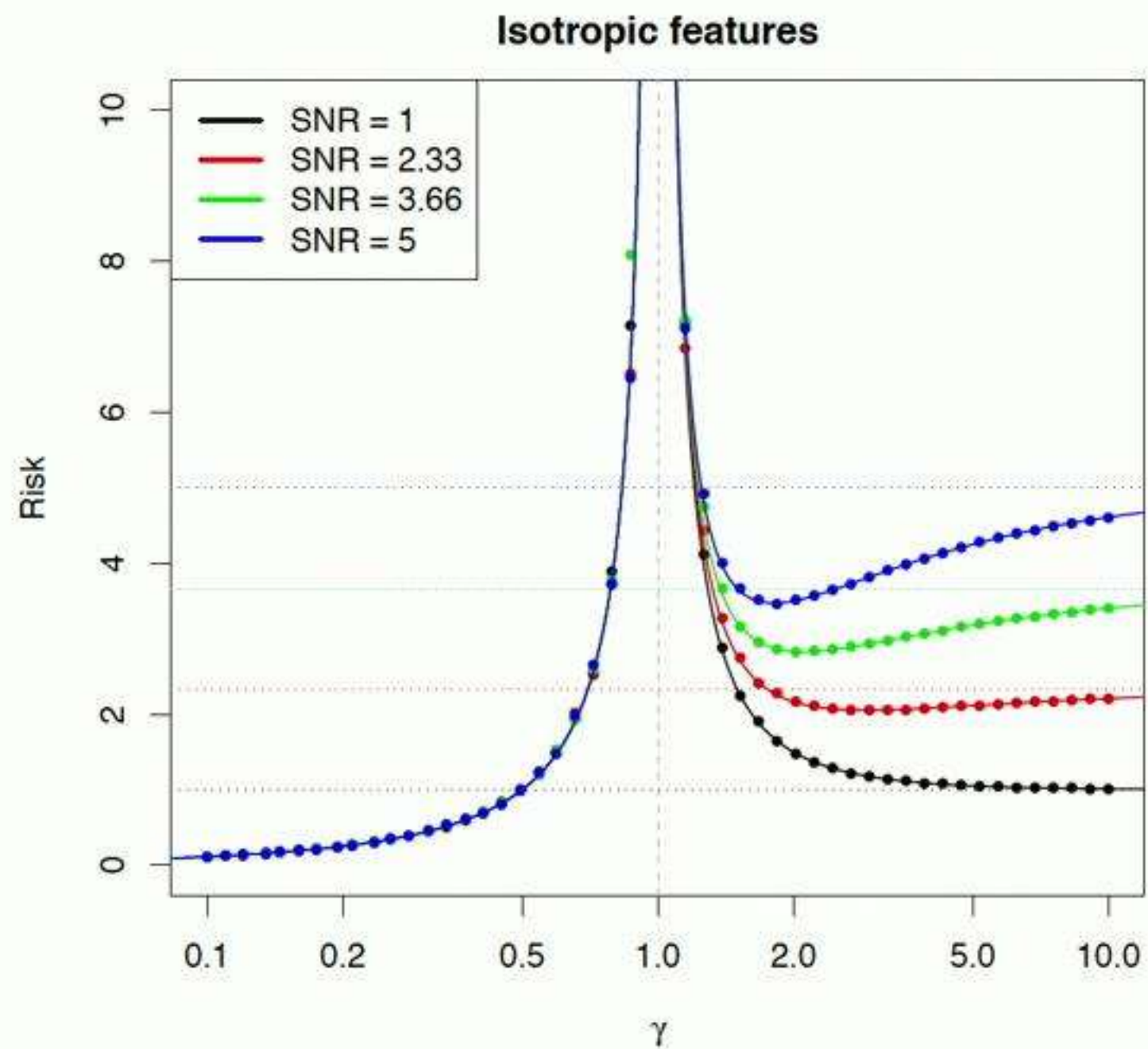
$$R(\beta_0) \rightarrow \begin{cases} \tau^2 \frac{\gamma}{1-\gamma} & \text{for } \gamma < 1, \\ \tau^2 \frac{1}{\gamma-1} + r^2 \left(1 - \frac{1}{\gamma}\right) & \text{for } \gamma > 1. \end{cases}$$

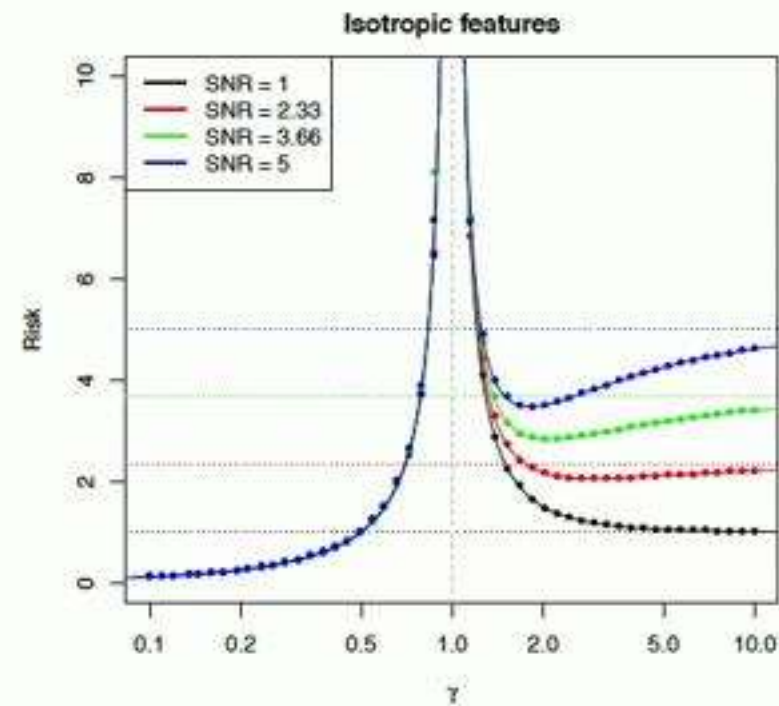
Further:

- ▶ If $\Sigma \neq \mathbf{I}_p$. . .
- ▶ If $z_i = \Sigma^{1/2} x_i$, x_i with iid components . . .

[Belkin, Hsu, Xu, 2019; Bartlett, Long, Lugosi, Tsigler, 2019]

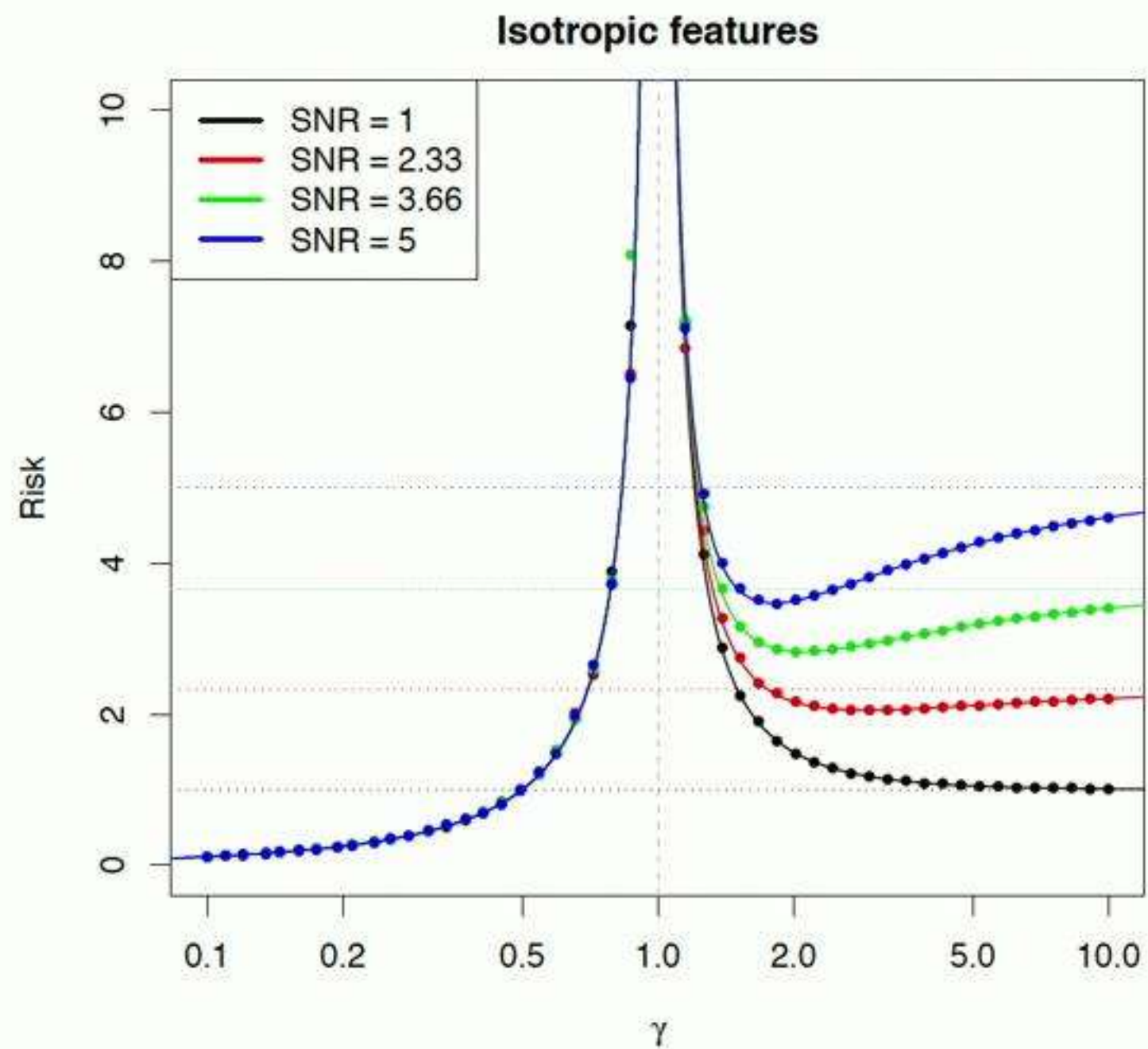
$$\Sigma = I_p, \text{ SNR} = \|\beta_0\|_2^2 / \tau^2$$





- ✓ Peak at the interpolation threshold
- ✗ Global minimum in the overparametrized regime
- ✗ Monotone decreasing in the overparametrized regime
- ✓ Vanishing (explicit) regularization
 - ▶ Dimension = Number of parameters

$$\Sigma = I_p, \text{ SNR} = \|\beta_0\|_2^2 / \tau^2$$



Precise asymptotics

Theorem (Hastie, M., Rosset, Tibshirani, 2019)

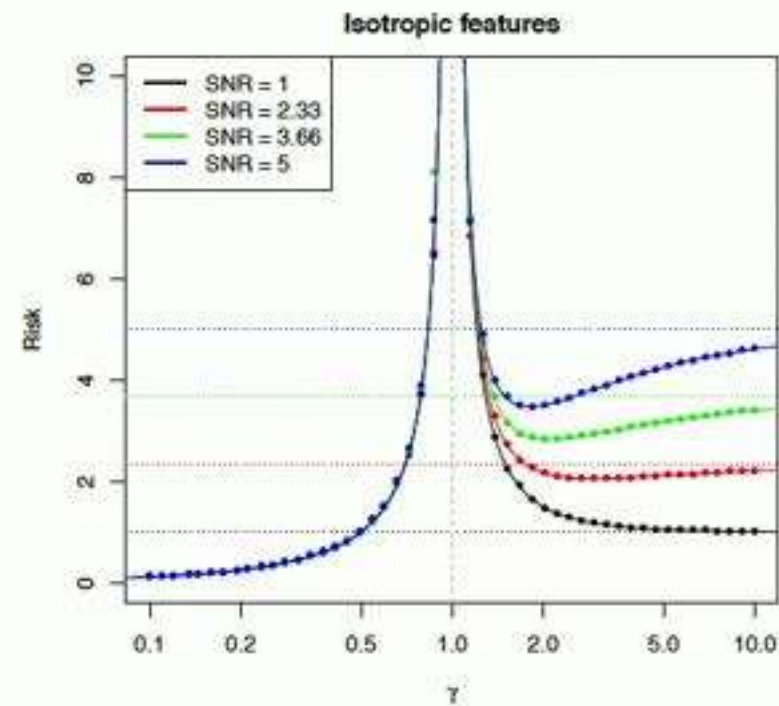
If $p/n \rightarrow \gamma$, $\Sigma = \mathbf{I}_p$, $\|\beta_0\|_2 \rightarrow r^2$, then

$$R(\beta_0) \rightarrow \begin{cases} \tau^2 \frac{\gamma}{1-\gamma} & \text{for } \gamma < 1, \\ \tau^2 \frac{1}{\gamma-1} + r^2 \left(1 - \frac{1}{\gamma}\right) & \text{for } \gamma > 1. \end{cases}$$

Further:

- ▶ If $\Sigma \neq \mathbf{I}_p$. . .
- ▶ If $z_i = \Sigma^{1/2} x_i$, x_i with iid components . . .

[Belkin, Hsu, Xu, 2019; Bartlett, Long, Lugosi, Tsigler, 2019]



- ✓ Peak at the interpolation threshold
- ✗ Global minimum in the overparametrized regime
- ✗ Monotone decreasing in the overparametrized regime
- ✓ Vanishing (explicit) regularization
 - ▶ Dimension = Number of parameters

Misspecified model

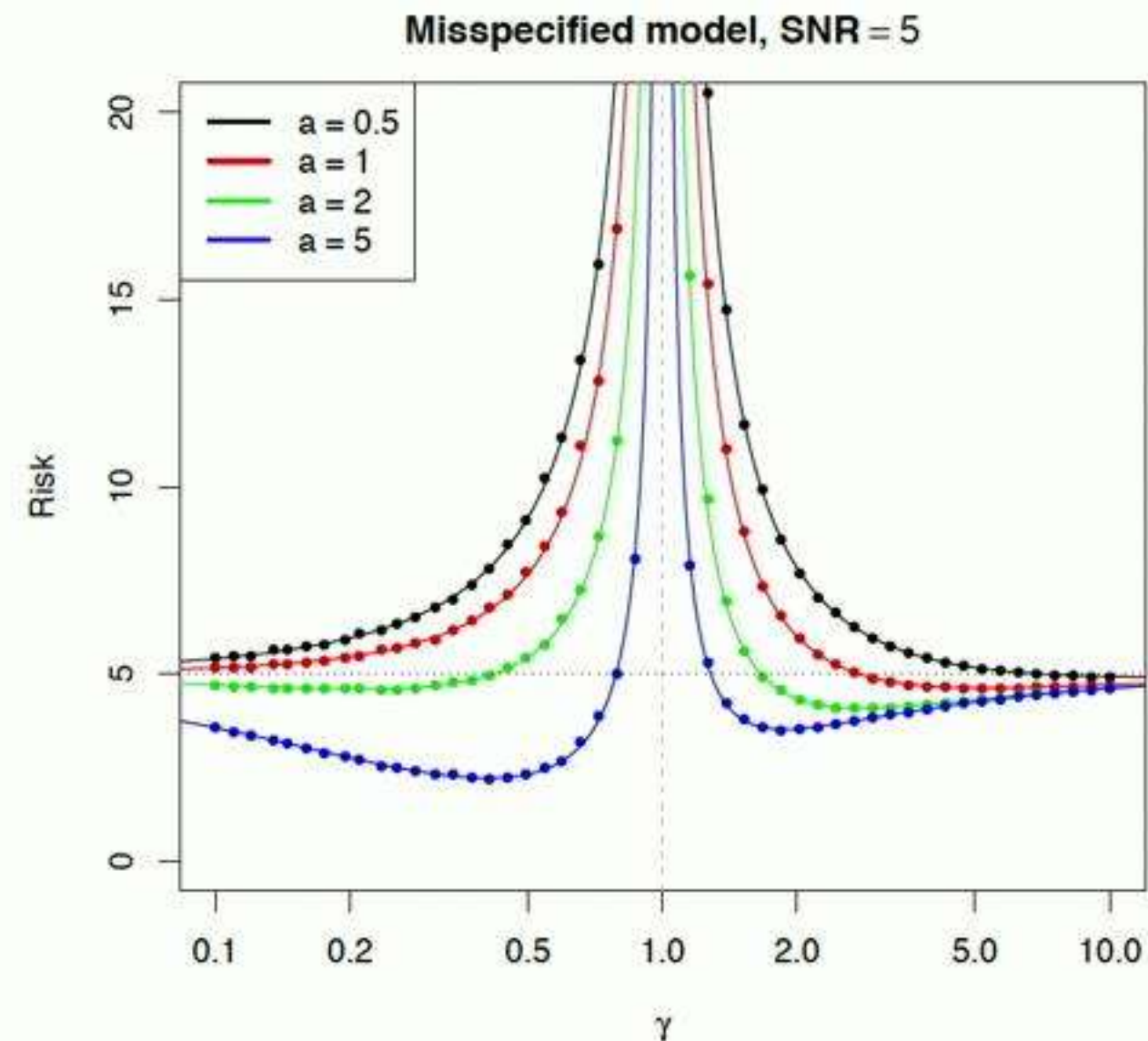


$$y_i = \langle \beta_0, z_i \rangle + \langle \tilde{\beta}_0, \tilde{z}_i \rangle + \varepsilon_i, \quad \varepsilon_i \sim \mathcal{N}(0, \tau^2)$$

- ▶ ‘Ridgeless regression’

$$\hat{\beta}(\lambda) \equiv \arg \min_{\beta \in \mathbb{R}^p} \left\{ \|y - Z\beta\|_2^2 + \lambda \|\beta\|_2^2 \right\},$$

Misspecified Gaussian model



- ✓ Global minimum in the overparametrized regime
- ✗ Monotone decreasing in the overparametrized regime
- Dimension $\neq$ Number of parameters

Nonlinear features: Approximation in high dimension

Nonparametric regression

Setting

- ▶ $\{(y_i, x_i)\}_{i \leq n}$ i.i.d. $x_i \sim \text{Unif}(\mathbb{S}^{d-1}(\sqrt{d}))$,

$$y_i = f_\star(x_i) + \varepsilon_i \quad \varepsilon_i \sim \mathcal{N}(0, \tau^2)$$

- ▶ Random features $w_i \sim \text{Unif}(\mathbb{S}^{d-1}(1))$,

$$\mathcal{F}_{\text{RF}}(W) \equiv \left\{ f(x) = \sum_{i=1}^N a_i \sigma(\langle w_i, x \rangle) : a_i \in \mathbb{R} \forall i \leq N \right\}.$$

- ▶ Ridge

$$\hat{a}(\lambda) = \operatorname{argmin}_{a \in \mathbb{R}^N} \left\{ \hat{\mathbb{E}}_n \left[\left(y - \sum_{i=1}^N a_i \sigma(\langle w_i, x \rangle) \right)^2 \right] + \frac{N\lambda}{d} \|a\|_2^2 \right\}.$$

Setting

- ▶ Test error

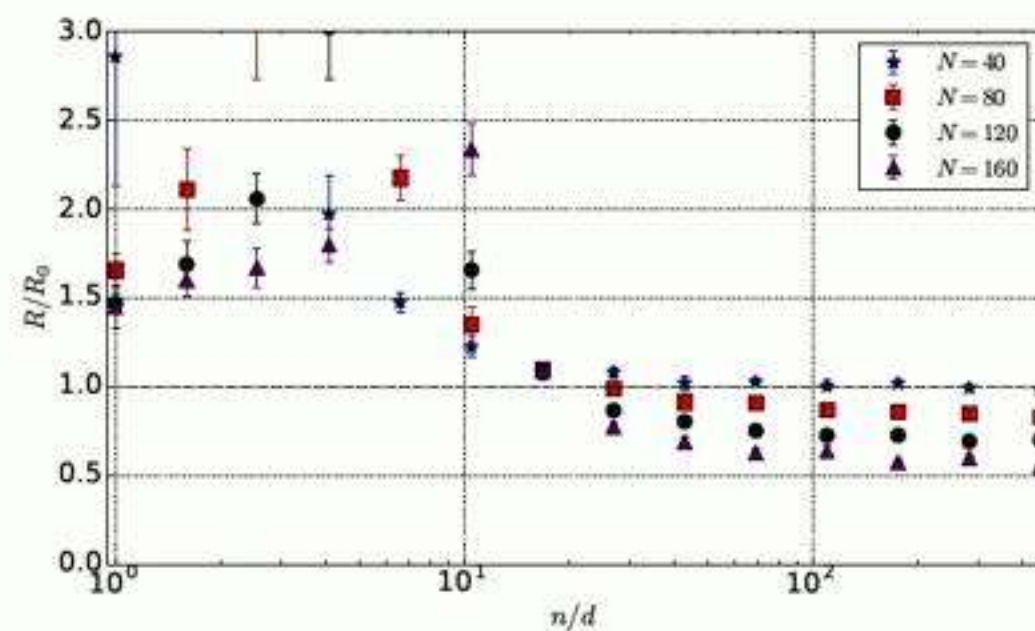
$$R_{\text{RF}}(\hat{f}; f) \equiv \mathbb{E}\{(\hat{f}(x) - f_{\star}(x))^2\}.$$

- ▶ First question

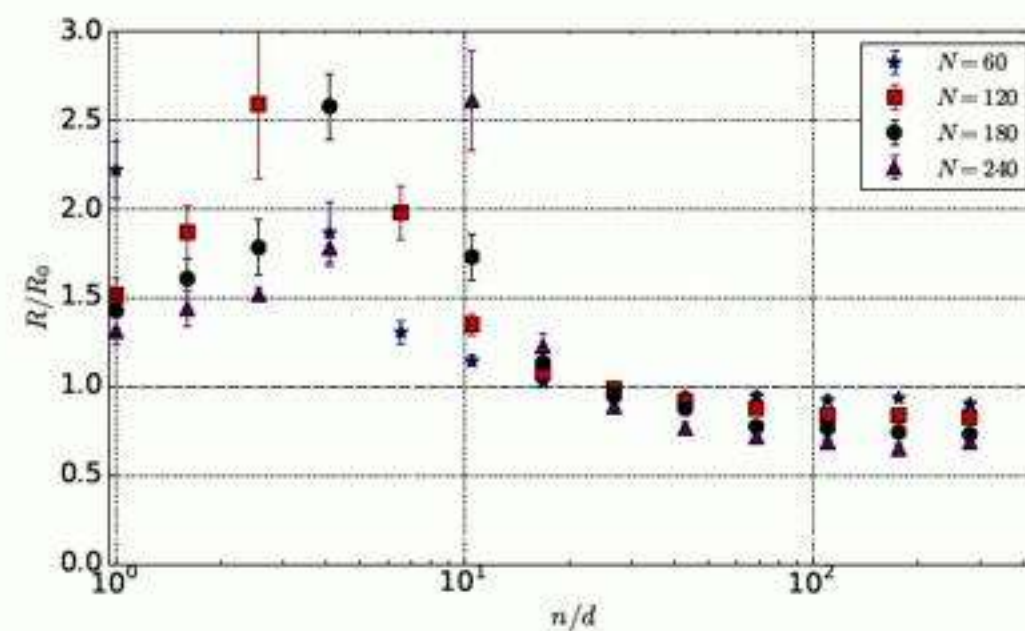
What can I fit with infinite data $n = \infty$?

An experiment

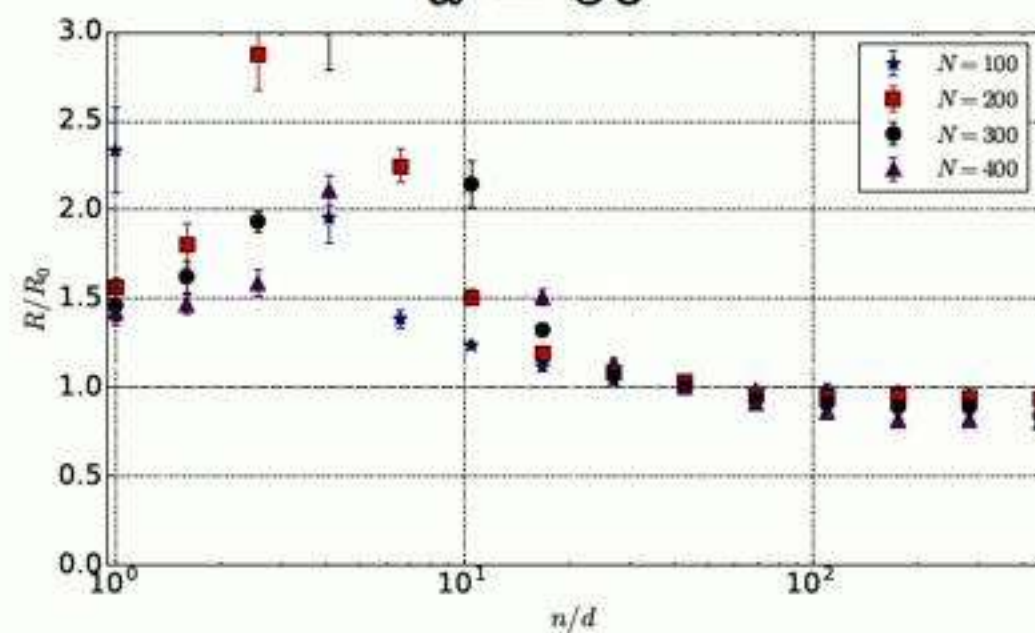
$d = 20$



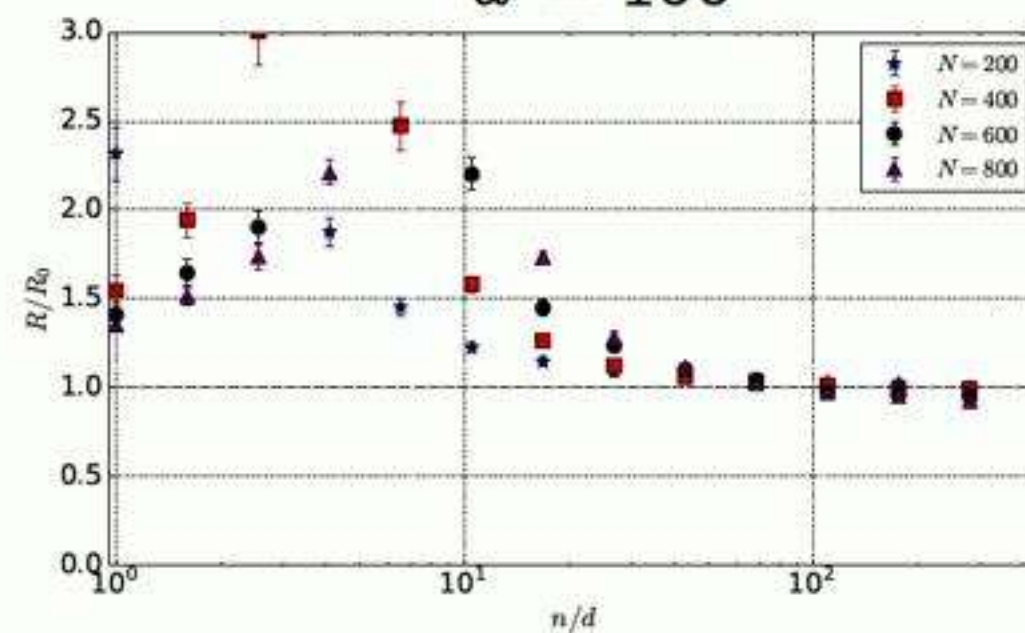
$d = 30$



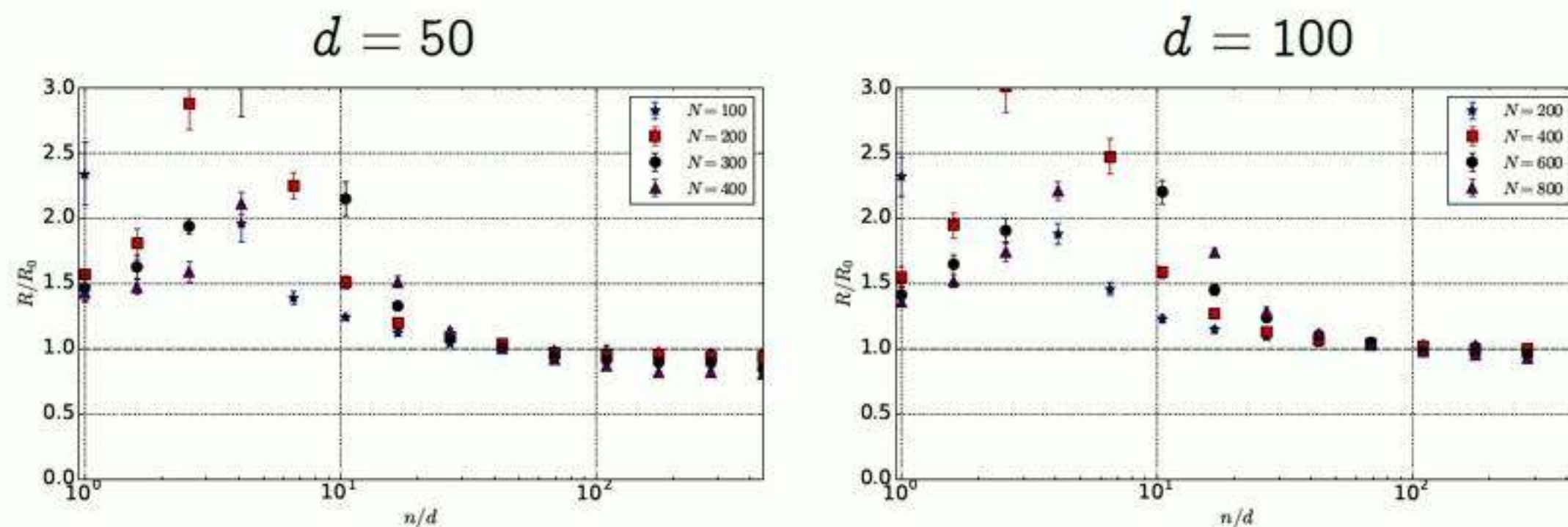
$d = 50$



$d = 100$



An experiment



- ▶ RF no better than trivial predictor $f_0(x) = 0$
- ▶ A 'difficult' target $f_\star$?

$$f_{\star,2}(x) = \sum_{i=1}^{d/2} x_i^2 - \sum_{i=d/2+1}^d x_i^2$$

A theorem $(P_{\leq \ell}$ projects on low degree)

Theorem (Ghorbani, Mei, Misiakiewicz, M. 2019)

Assume $N \leq d^{\ell+1-\delta}$ for a fixed integer ℓ and some $\delta > 0$.

Then, with high probability:

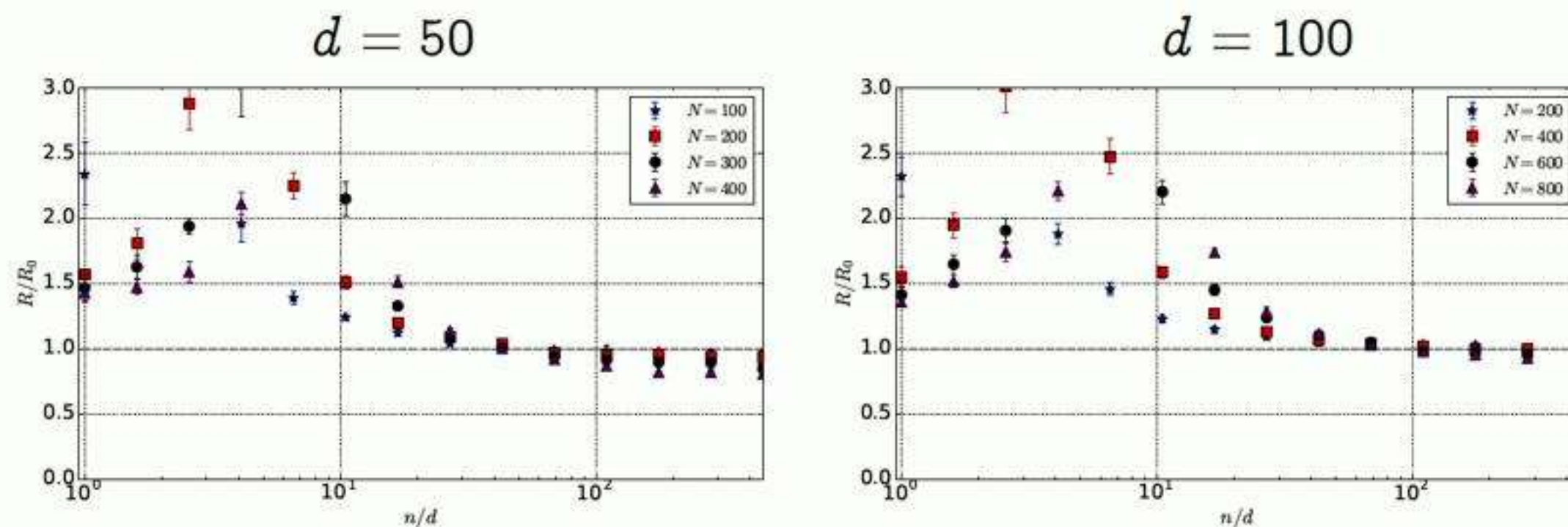
$$R_{\text{RF}}(f_{\star}) = R_{\text{RF}}(P_{\leq \ell} f_{\star}) + \|P_{> \ell} f_{\star}\|_{L^2}^2 + o(\|f_{\star}\|_{L^2}^2).$$

Example: If # neurons $N \ll d^2$, then

$$\begin{aligned} R_{\text{RF}}(f_{\star}, W) &\gtrsim R_{\text{lin. reg.}}(f_{\star}) \\ &\equiv \min_{b, \beta} \mathbb{E}\{(f_{\star}(x) - b - \langle \beta, x \rangle)\}^2. \end{aligned}$$

[Bach 2017;...]

An experiment



- ▶ RF no better than trivial predictor $f_0(x) = 0$
- ▶ A 'difficult' target $f_\star$?

$$f_{\star,2}(x) = \sum_{i=1}^{d/2} x_i^2 - \sum_{i=d/2+1}^d x_i^2$$

A theorem $(P_{\leq \ell}$ projects on low degree)

Theorem (Ghorbani, Mei, Misiakiewicz, M. 2019)

Assume $N \leq d^{\ell+1-\delta}$ for a fixed integer ℓ and some $\delta > 0$.

Then, with high probability:

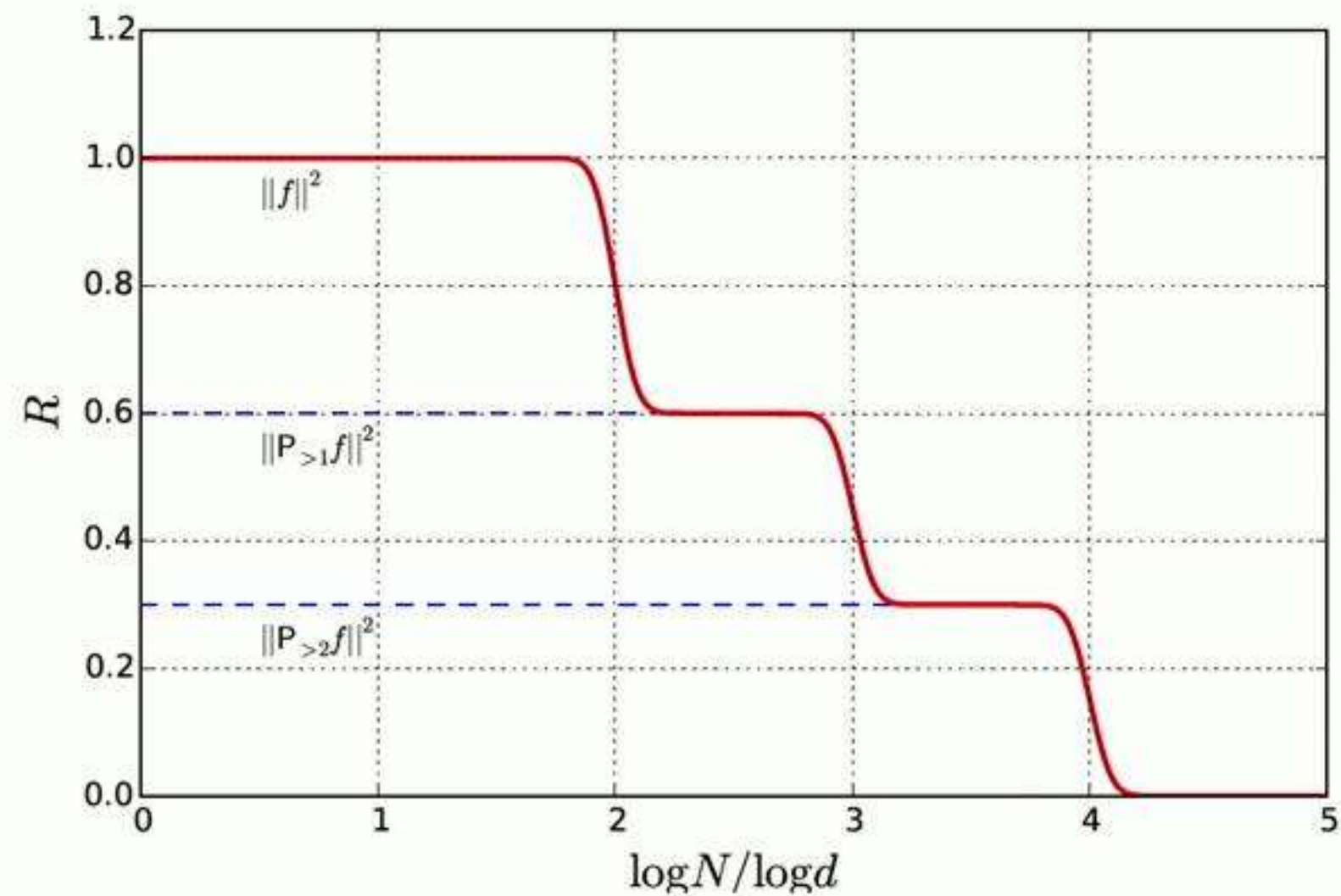
$$R_{\text{RF}}(f_{\star}) = R_{\text{RF}}(P_{\leq \ell} f_{\star}) + \|P_{> \ell} f_{\star}\|_{L^2}^2 + o(\|f_{\star}\|_{L^2}^2).$$

Example: If # neurons $N \ll d^2$, then

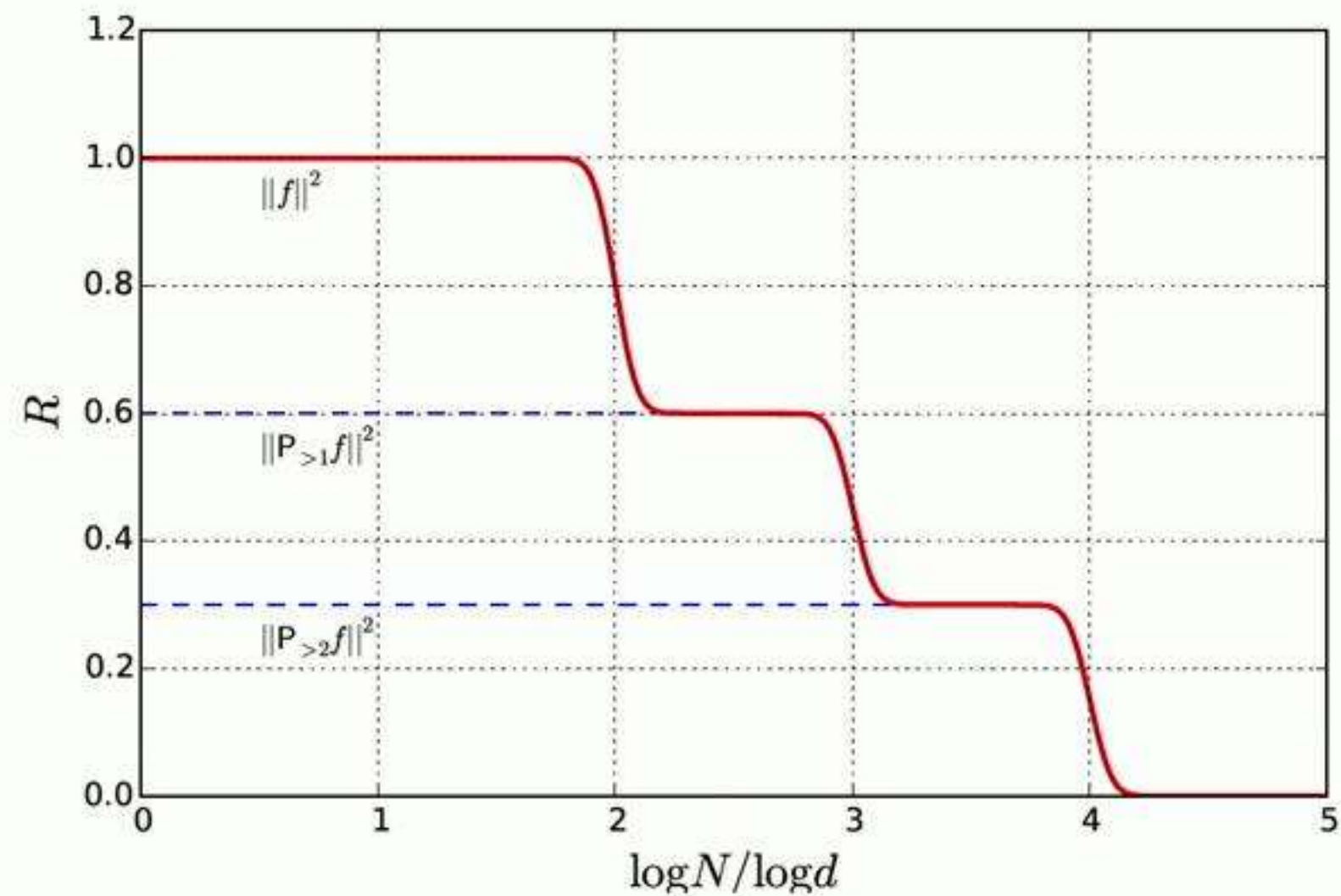
$$\begin{aligned} R_{\text{RF}}(f_{\star}, W) &\gtrsim R_{\text{lin. reg.}}(f_{\star}) \\ &\equiv \min_{b, \beta} \mathbb{E}\{ (f_{\star}(x) - b - \langle \beta, x \rangle) \}. \end{aligned}$$

[Bach 2017;...]

Similar theorem for NT



What happens at finite n ?



What happens at finite n ?

Nonlinear features: Generalization curve

Setting

- ▶ Random features ridge regression (as above.)
- ▶ Proportional regime $n, N, d \rightarrow \infty$,

$$\frac{N}{d} \rightarrow \psi_1, \quad \frac{n}{d} \rightarrow \psi_2.$$

- ▶ $f_\star(x) = \langle \beta_0, x \rangle$ (Higher order terms $\approx$ noise)

Precise asymptotics

Theorem

Assume $f_\star(x) = \langle \beta_0, x \rangle$ and define (for $G \sim N(0, 1)$)

$$b_\star^2 = \mathbb{E}[\sigma(G)^2] - \mathbb{E}[\sigma(G)]^2 - b_1^2, \quad \zeta \equiv \frac{b_1^2}{b_\star^2}.$$

Then, for any $\lambda > 0$, we have

$$R_{\text{RF}}(\hat{f}_\lambda, f_\star) = \|\beta_0\|_2^2 \mathcal{B}(\zeta, \psi_1, \psi_2, \lambda/\bar{b}_\star^2) + \tau^2 \mathcal{V}(\zeta, \psi_1, \psi_2, \lambda/\bar{b}_\star^2) + o_d(1),$$

where $\mathcal{B}(\zeta, \psi_1, \psi_2, \bar{\lambda})$, $\mathcal{V}(\zeta, \psi_1, \psi_2, \bar{\lambda})$ are explicitly given below.

Explicit formulae

Let $(\nu_1(\xi), \nu_2(\xi))$ be the unique solution of

$$\begin{aligned}\nu_1 &= \psi_1 \left(-\xi - \nu_2 - \frac{\zeta^2 \nu_2}{1 - \zeta^2 \nu_1 \nu_2} \right)^{-1}, \\ \nu_2 &= \psi_2 \left(-\xi - \nu_1 - \frac{\zeta^2 \nu_1}{1 - \zeta^2 \nu_1 \nu_2} \right)^{-1};\end{aligned}$$

Let

$$\chi \equiv \nu_1(i(\psi_1 \psi_2 \bar{\lambda})^{1/2}) \cdot \nu_2(i(\psi_1 \psi_2 \bar{\lambda})^{1/2}),$$

and

$$\begin{aligned}\mathcal{E}_0(\zeta, \psi_1, \psi_2, \bar{\lambda}) &\equiv -\chi^5 \zeta^6 + 3\chi^4 \zeta^4 + (\psi_1 \psi_2 - \psi_2 - \psi_1 + 1)\chi^3 \zeta^6 - 2\chi^3 \zeta^4 - 3\chi^3 \zeta^2 \\ &\quad + (\psi_1 + \psi_2 - 3\psi_1 \psi_2 + 1)\chi^2 \zeta^4 + 2\chi^2 \zeta^2 + \chi^2 + 3\psi_1 \psi_2 \chi \zeta^2 - \psi_1 \psi_2, \\ \mathcal{E}_1(\zeta, \psi_1, \psi_2, \bar{\lambda}) &\equiv \psi_2 \chi^3 \zeta^4 - \psi_2 \chi^2 \zeta^2 + \psi_1 \psi_2 \chi \zeta^2 - \psi_1 \psi_2, \\ \mathcal{E}_2(\zeta, \psi_1, \psi_2, \bar{\lambda}) &\equiv \chi^5 \zeta^6 - 3\chi^4 \zeta^4 + (\psi_1 - 1)\chi^3 \zeta^6 + 2\chi^3 \zeta^4 + 3\chi^3 \zeta^2 + (-\psi_1 - 1)\chi^2 \zeta^4 - 2\chi^2 \zeta^2 - \chi^2.\end{aligned}$$

We then have

$$\mathcal{B}(\zeta, \psi_1, \psi_2, \bar{\lambda}) \equiv \frac{\mathcal{E}_1(\zeta, \psi_1, \psi_2, \bar{\lambda})}{\mathcal{E}_0(\zeta, \psi_1, \psi_2, \bar{\lambda})}, \quad \mathcal{V}(\zeta, \psi_1, \psi_2, \bar{\lambda}) \equiv \frac{\mathcal{E}_2(\zeta, \psi_1, \psi_2, \bar{\lambda})}{\mathcal{E}_0(\zeta, \psi_1, \psi_2, \bar{\lambda})}.$$

Random matrix theory for kernel inner product random matrices

[Cheng, Singer, 2013; Do, Vu 2013; Fan, M. 2019; ...]

Explicit formulae

Let $(\nu_1(\xi), \nu_2(\xi))$ be the unique solution of

$$\begin{aligned}\nu_1 &= \psi_1 \left(-\xi - \nu_2 - \frac{\zeta^2 \nu_2}{1 - \zeta^2 \nu_1 \nu_2} \right)^{-1}, \\ \nu_2 &= \psi_2 \left(-\xi - \nu_1 - \frac{\zeta^2 \nu_1}{1 - \zeta^2 \nu_1 \nu_2} \right)^{-1};\end{aligned}$$

Let

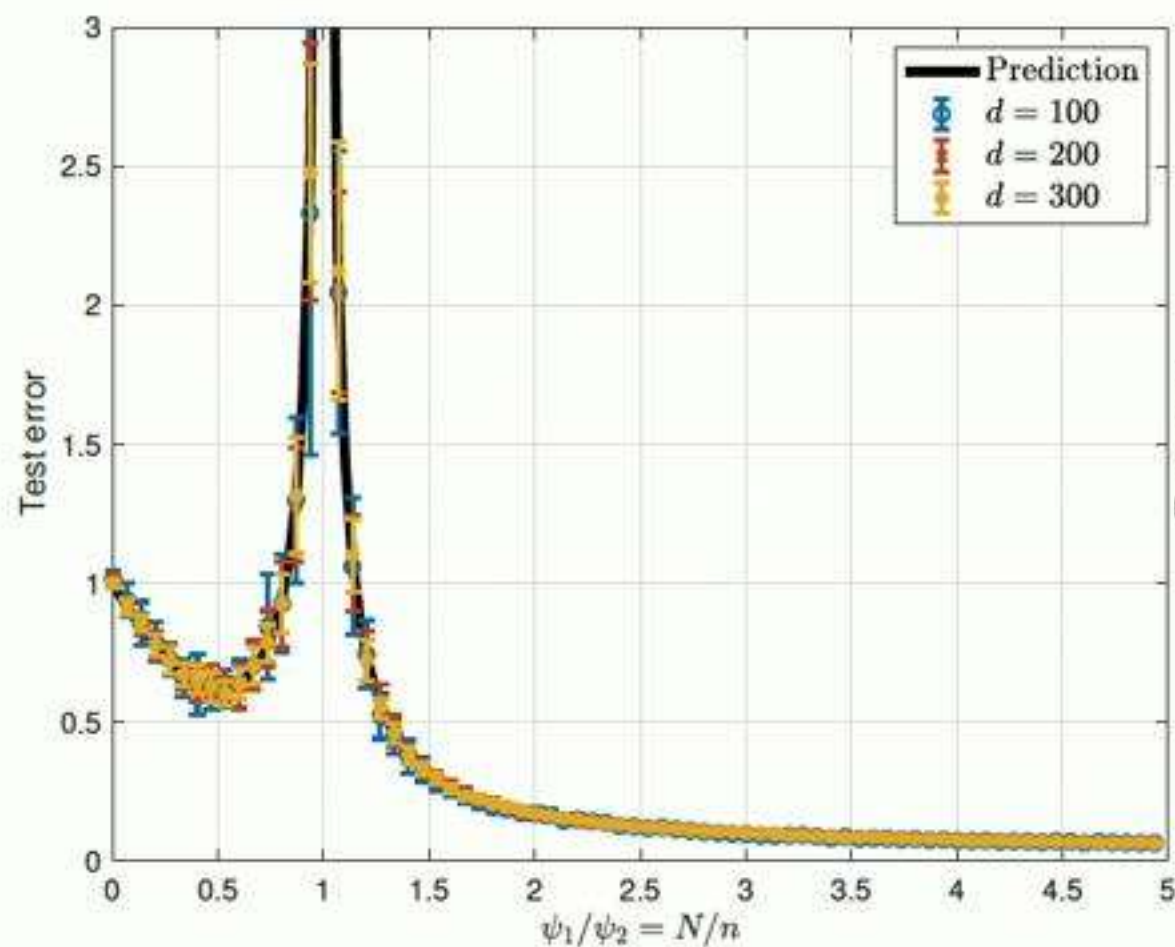
$$\chi \equiv \nu_1(i(\psi_1 \psi_2 \bar{\lambda})^{1/2}) \cdot \nu_2(i(\psi_1 \psi_2 \bar{\lambda})^{1/2}),$$

and

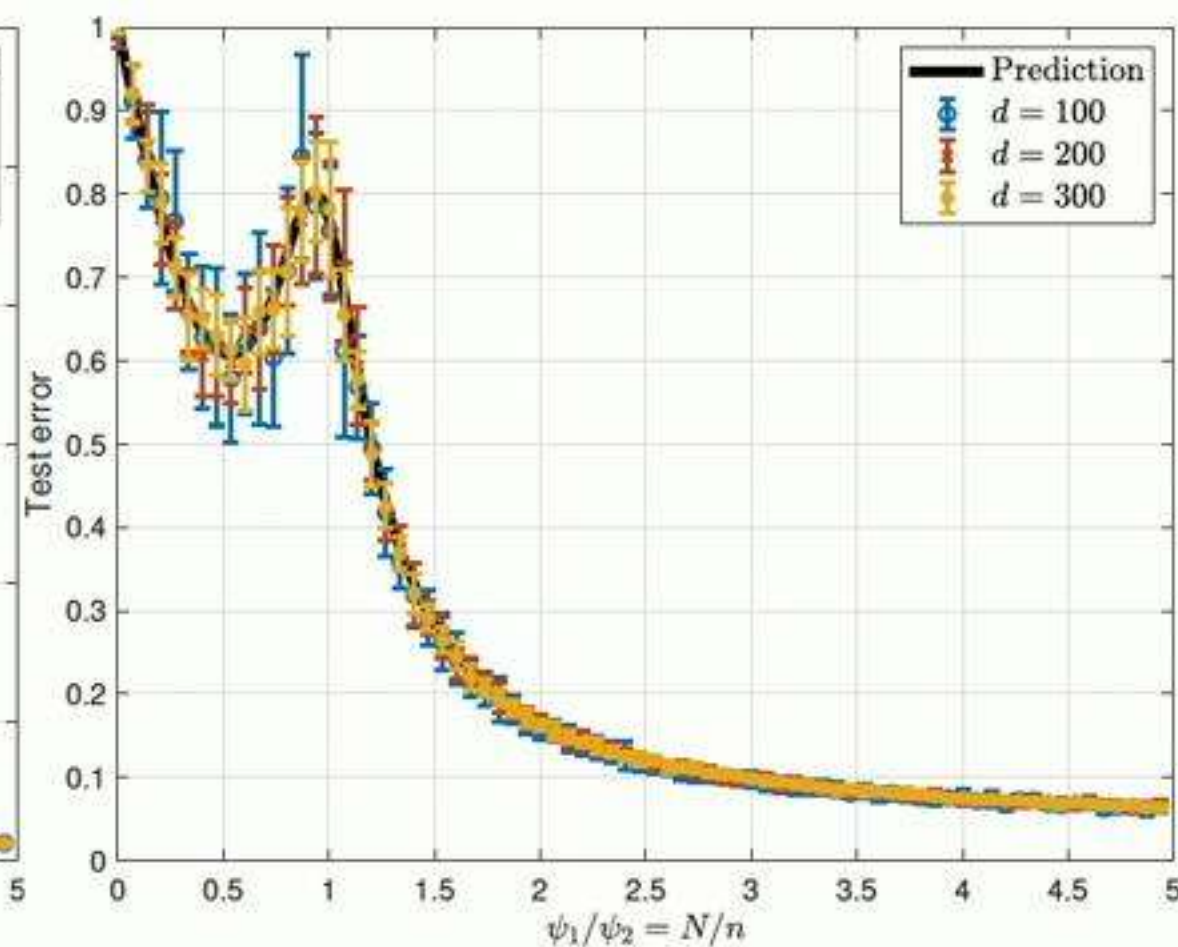
$$\begin{aligned}\mathcal{E}_0(\zeta, \psi_1, \psi_2, \bar{\lambda}) &\equiv -\chi^5 \zeta^6 + 3\chi^4 \zeta^4 + (\psi_1 \psi_2 - \psi_2 - \psi_1 + 1)\chi^3 \zeta^6 - 2\chi^3 \zeta^4 - 3\chi^3 \zeta^2 \\ &\quad + (\psi_1 + \psi_2 - 3\psi_1 \psi_2 + 1)\chi^2 \zeta^4 + 2\chi^2 \zeta^2 + \chi^2 + 3\psi_1 \psi_2 \chi \zeta^2 - \psi_1 \psi_2, \\ \mathcal{E}_1(\zeta, \psi_1, \psi_2, \bar{\lambda}) &\equiv \psi_2 \chi^3 \zeta^4 - \psi_2 \chi^2 \zeta^2 + \psi_1 \psi_2 \chi \zeta^2 - \psi_1 \psi_2, \\ \mathcal{E}_2(\zeta, \psi_1, \psi_2, \bar{\lambda}) &\equiv \chi^5 \zeta^6 - 3\chi^4 \zeta^4 + (\psi_1 - 1)\chi^3 \zeta^6 + 2\chi^3 \zeta^4 + 3\chi^3 \zeta^2 + (-\psi_1 - 1)\chi^2 \zeta^4 - 2\chi^2 \zeta^2 - \chi^2.\end{aligned}$$

We then have

$$\mathcal{B}(\zeta, \psi_1, \psi_2, \bar{\lambda}) \equiv \frac{\mathcal{E}_1(\zeta, \psi_1, \psi_2, \bar{\lambda})}{\mathcal{E}_0(\zeta, \psi_1, \psi_2, \bar{\lambda})}, \quad \mathcal{V}(\zeta, \psi_1, \psi_2, \bar{\lambda}) \equiv \frac{\mathcal{E}_2(\zeta, \psi_1, \psi_2, \bar{\lambda})}{\mathcal{E}_0(\zeta, \psi_1, \psi_2, \bar{\lambda})}.$$

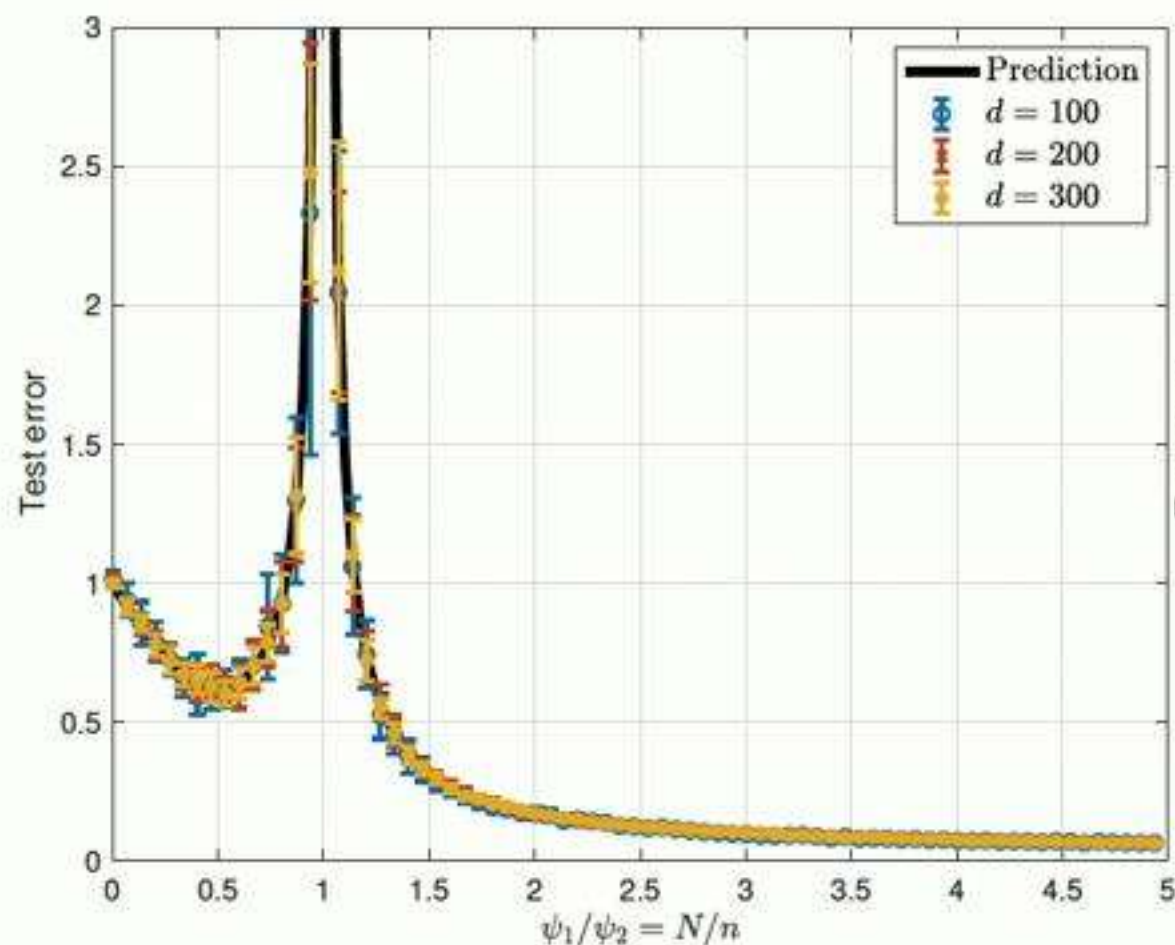


$\lambda = 0+$

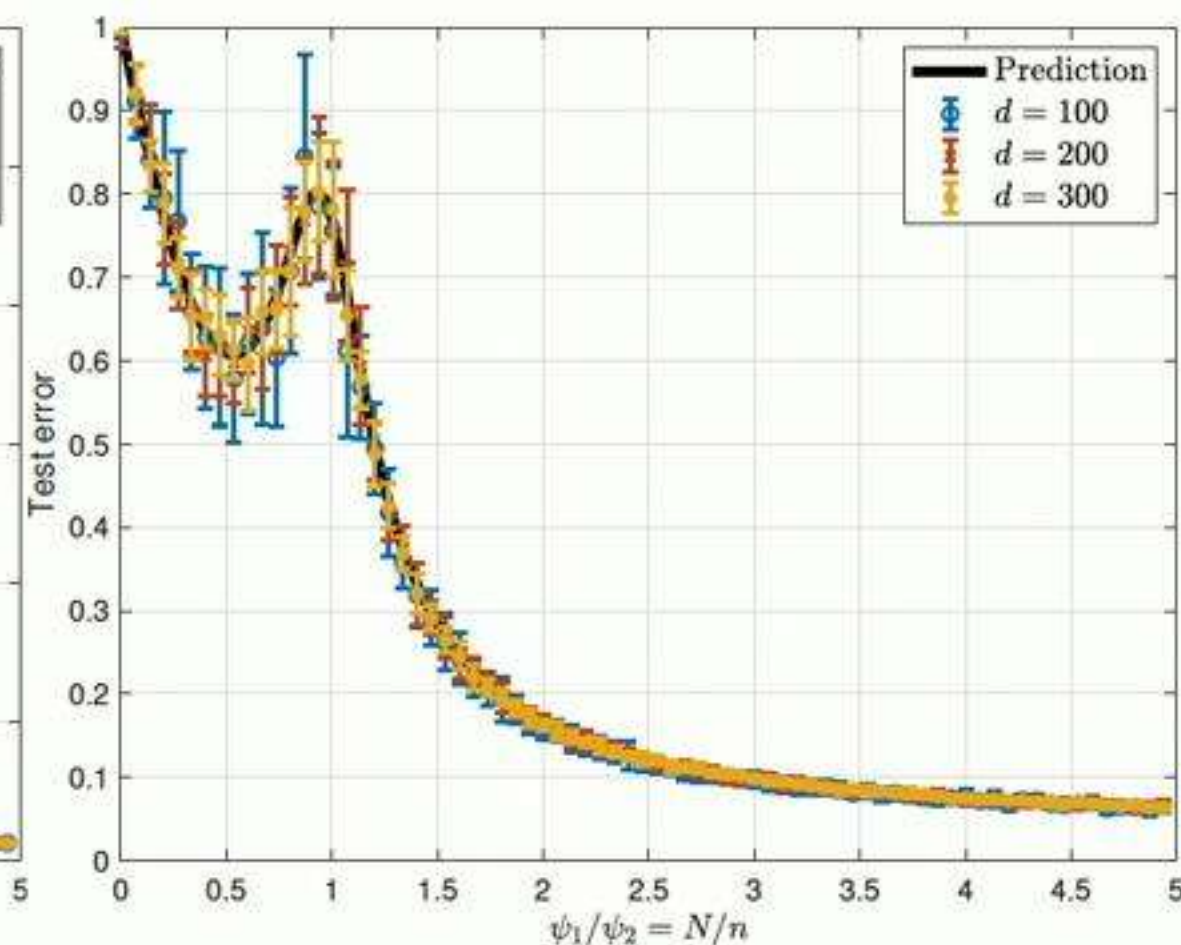


$\lambda > 0$

- ✓ Peak at the interpolation threshold
- ✓ Global minimum in the overparametrized regime
- ✓ Monotone decreasing in the overparametrized regime



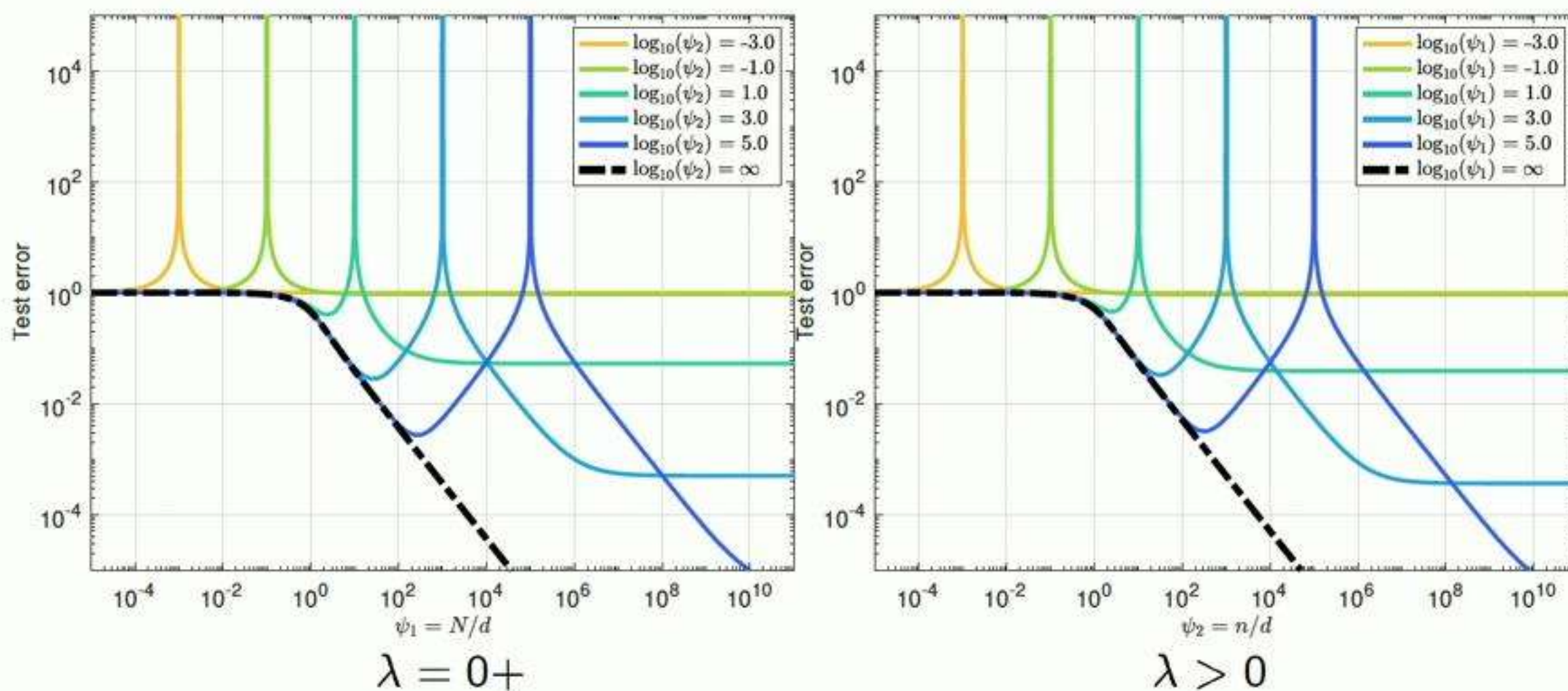
$\lambda = 0+$



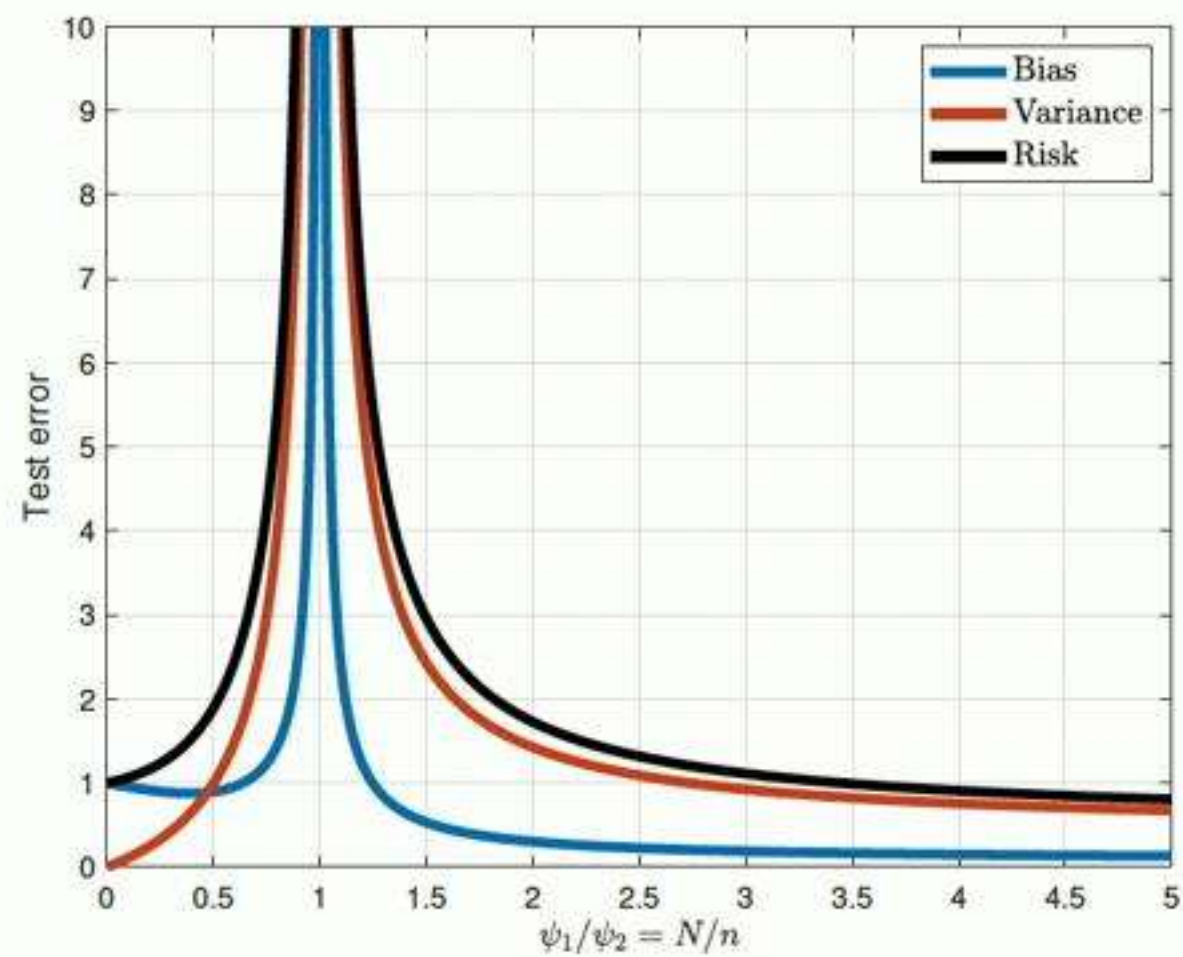
$\lambda > 0$

- ✓ Peak at the interpolation threshold
- ✓ Global minimum in the overparametrized regime
- ✓ Monotone decreasing in the overparametrized regime

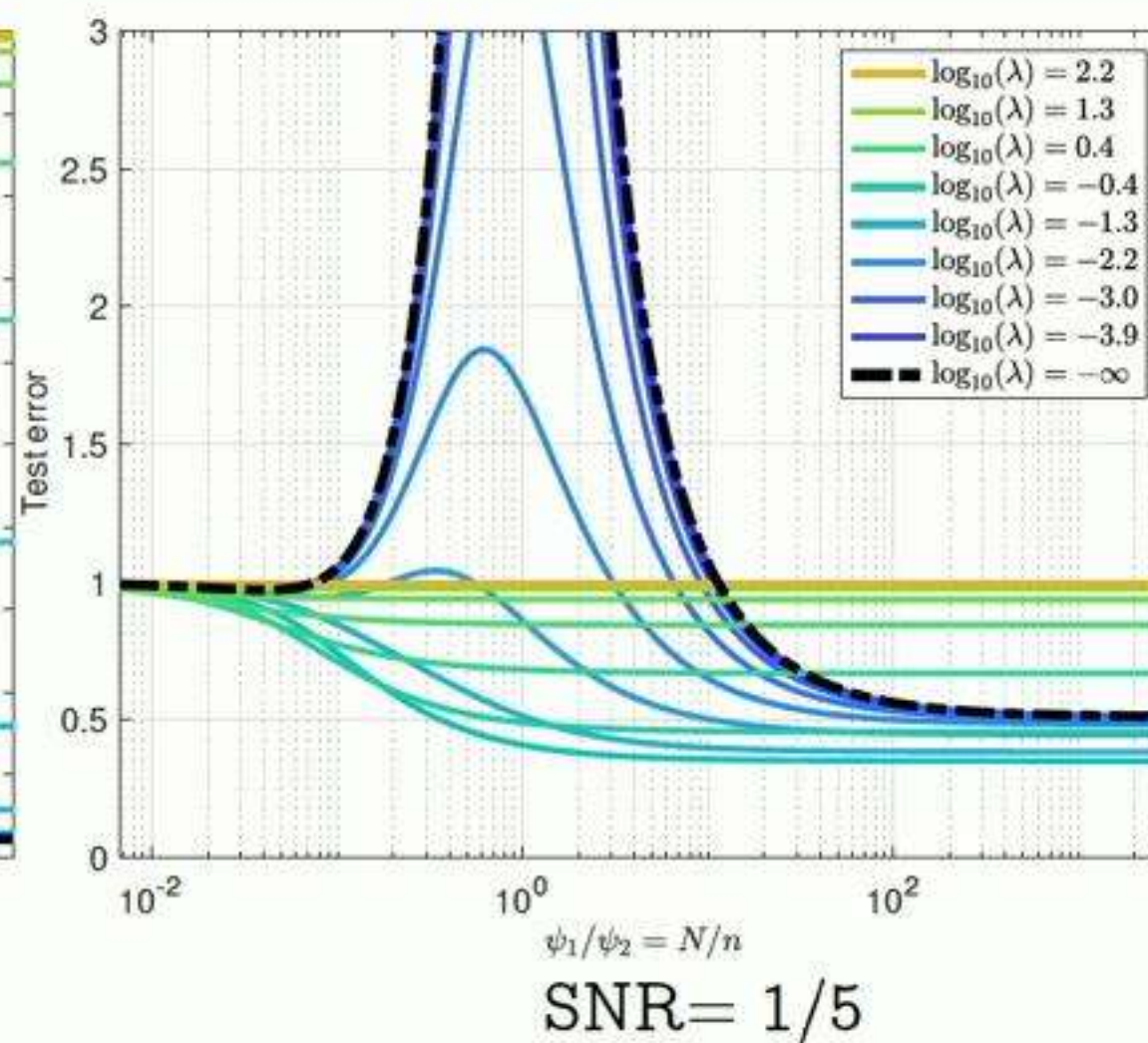
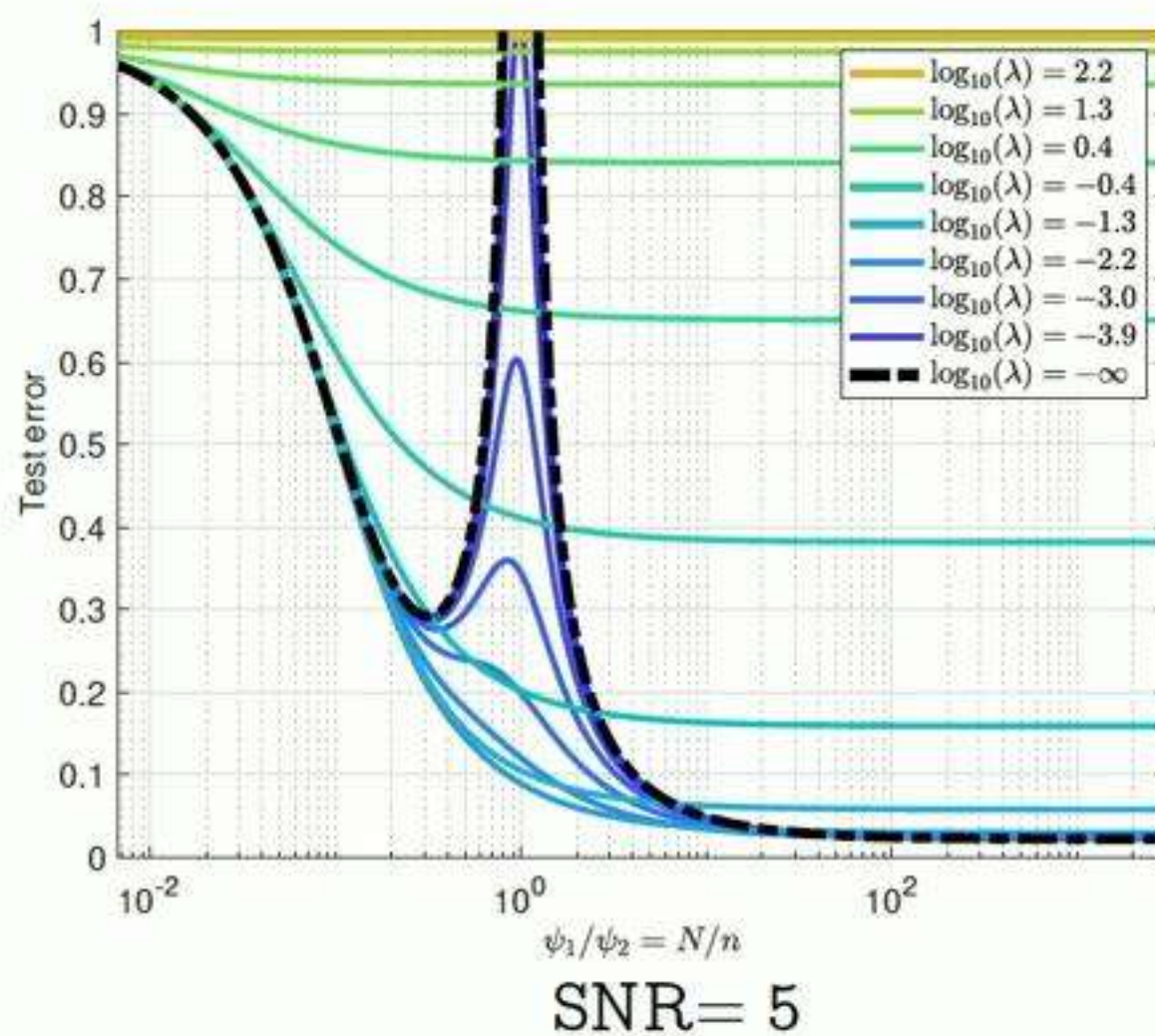
Insights



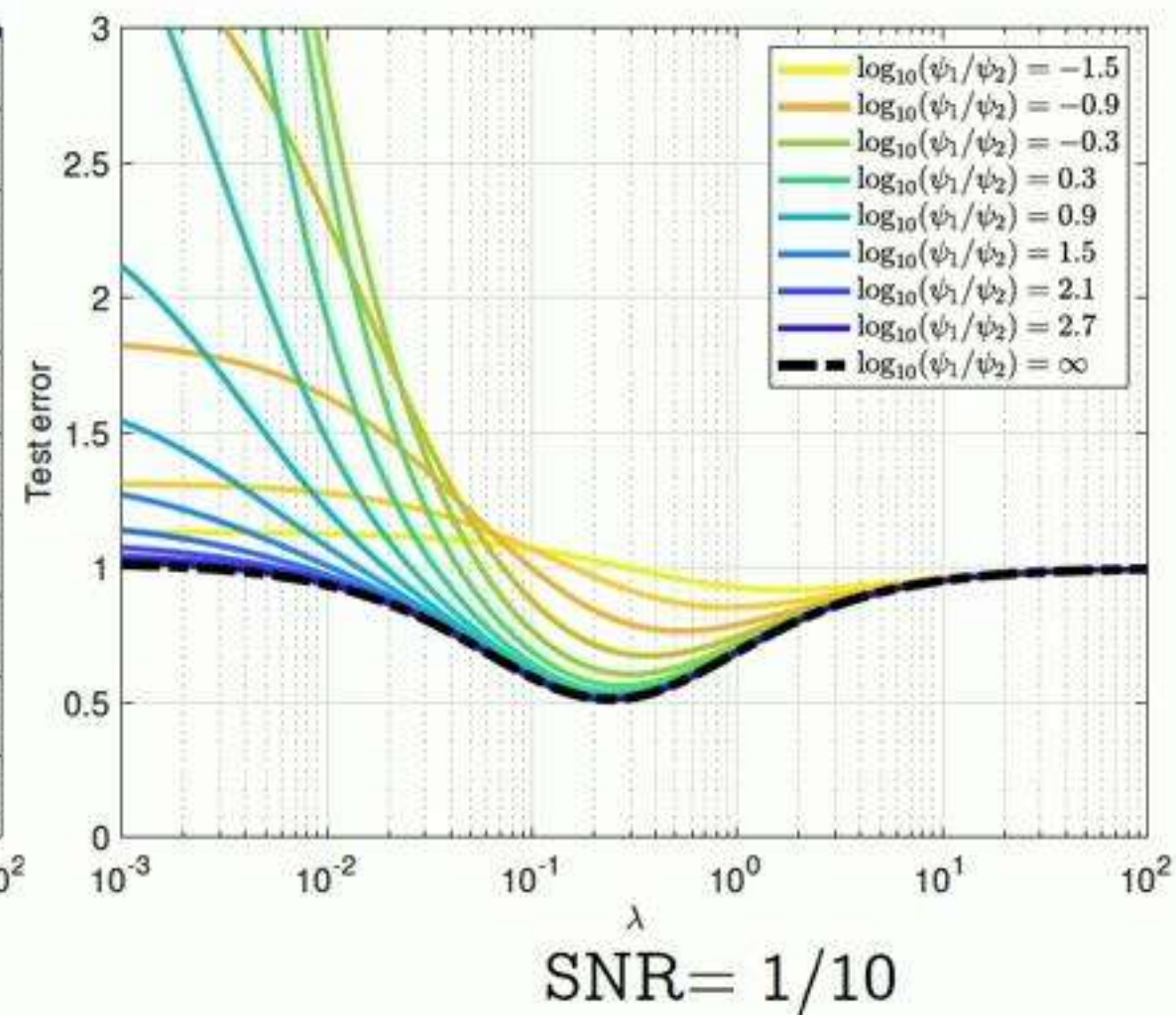
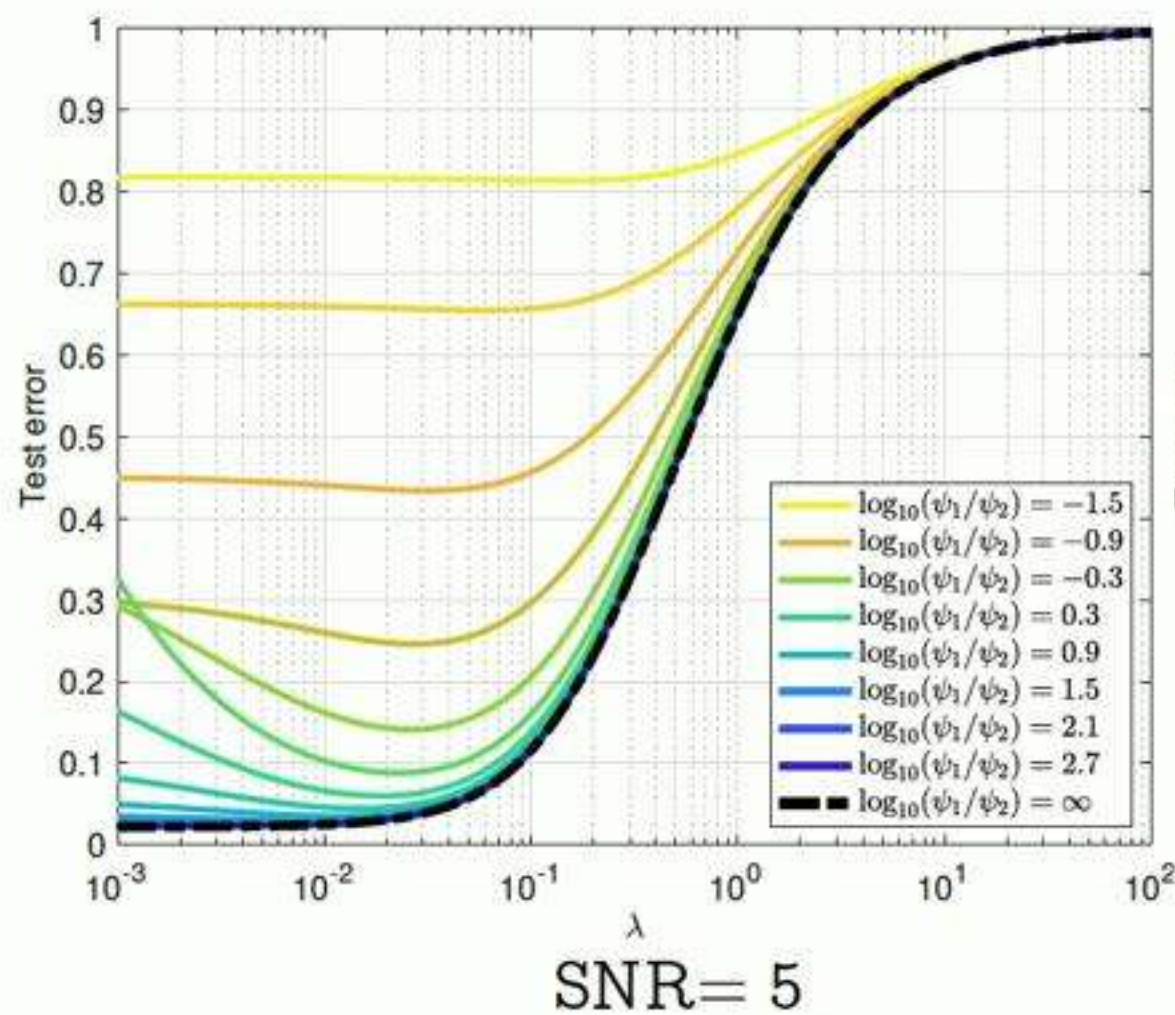
- Singularity at the interpolation threshold
- Minimum risk at extreme overparametrization $N/n \rightarrow \infty$.



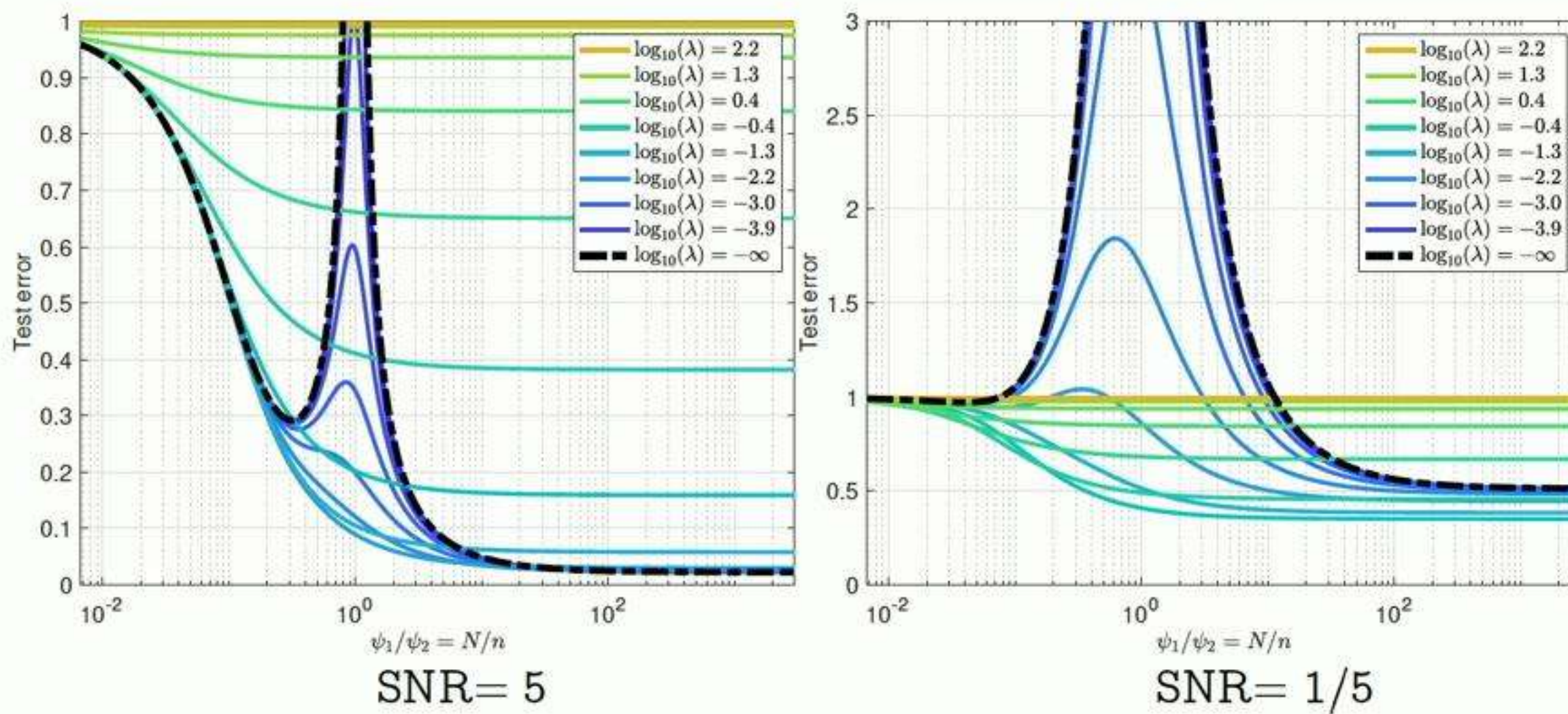
- Both bias and variance are singular



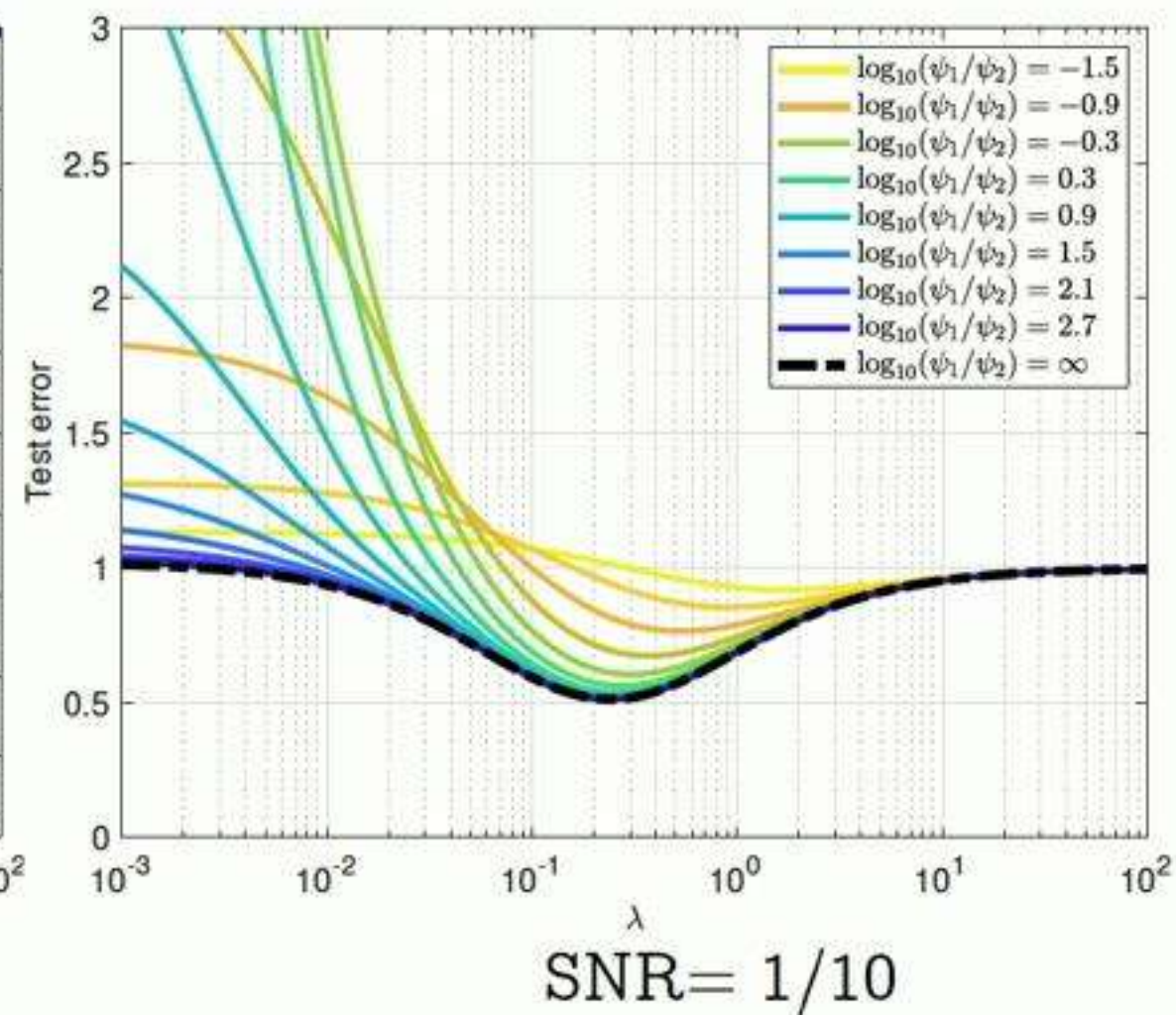
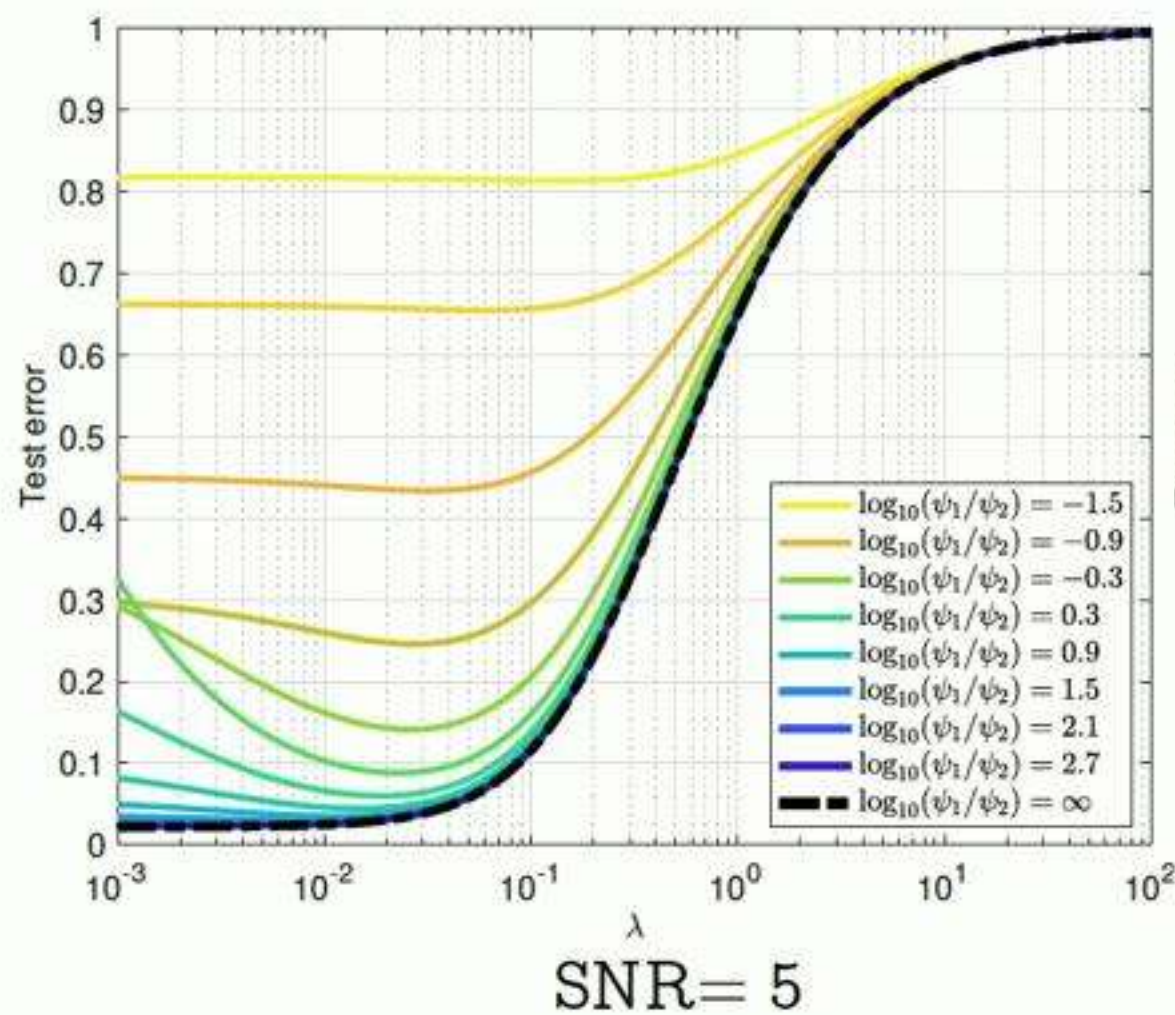
- Same at $\lambda > 0$ fixed: Minimum at $N/n \rightarrow \infty$.



- ▶ High SNR: Minimum at $\lambda = 0+$.
- ▶ Low SNR: Minimum at $\lambda > 0$.



- Same at $\lambda > 0$ fixed: Minimum at $N/n \rightarrow \infty$.



- ▶ High SNR: Minimum at $\lambda = 0+$.
- ▶ Low SNR: Minimum at $\lambda > 0$.

Conclusion

Overparametrization $\nRightarrow$ Poor generalization

- ▶ Mechanism #1: Early stopping/mean field
[Mei, M, Nguyen 2018; Chizat, Bach 2018]
- ▶ Mechanism #2: High dimension/interpolation
[This talk]
- ▶ Open questions:
 - ▶ Fully trained networks?
 - ▶ When is extreme overparametrization optimal?

Thanks!

Overparametrization $\nRightarrow$ Poor generalization

- ▶ Mechanism #1: Early stopping/mean field
[Mei, M, Nguyen 2018; Chizat, Bach 2018]
- ▶ Mechanism #2: High dimension/interpolation
[This talk]
- ▶ Open questions:
 - ▶ Fully trained networks?
 - ▶ When is extreme overparametrization optimal?

Thanks!

Explicit formulae

Let $(\nu_1(\xi), \nu_2(\xi))$ be the unique solution of

$$\begin{aligned}\nu_1 &= \psi_1 \left(-\xi - \nu_2 - \frac{\zeta^2 \nu_2}{1 - \zeta^2 \nu_1 \nu_2} \right)^{-1}, \\ \nu_2 &= \psi_2 \left(-\xi - \nu_1 - \frac{\zeta^2 \nu_1}{1 - \zeta^2 \nu_1 \nu_2} \right)^{-1};\end{aligned}$$

Let

$$\chi \equiv \nu_1(i(\psi_1 \psi_2 \bar{\lambda})^{1/2}) \cdot \nu_2(i(\psi_1 \psi_2 \bar{\lambda})^{1/2}),$$

and

$$\begin{aligned}\mathcal{E}_0(\zeta, \psi_1, \psi_2, \bar{\lambda}) &\equiv -\chi^5 \zeta^6 + 3\chi^4 \zeta^4 + (\psi_1 \psi_2 - \psi_2 - \psi_1 + 1)\chi^3 \zeta^6 - 2\chi^3 \zeta^4 - 3\chi^3 \zeta^2 \\ &\quad + (\psi_1 + \psi_2 - 3\psi_1 \psi_2 + 1)\chi^2 \zeta^4 + 2\chi^2 \zeta^2 + \chi^2 + 3\psi_1 \psi_2 \chi \zeta^2 - \psi_1 \psi_2, \\ \mathcal{E}_1(\zeta, \psi_1, \psi_2, \bar{\lambda}) &\equiv \psi_2 \chi^3 \zeta^4 - \psi_2 \chi^2 \zeta^2 + \psi_1 \psi_2 \chi \zeta^2 - \psi_1 \psi_2, \\ \mathcal{E}_2(\zeta, \psi_1, \psi_2, \bar{\lambda}) &\equiv \chi^5 \zeta^6 - 3\chi^4 \zeta^4 + (\psi_1 - 1)\chi^3 \zeta^6 + 2\chi^3 \zeta^4 + 3\chi^3 \zeta^2 + (-\psi_1 - 1)\chi^2 \zeta^4 - 2\chi^2 \zeta^2 - \chi^2.\end{aligned}$$

We then have

$$\mathcal{B}(\zeta, \psi_1, \psi_2, \bar{\lambda}) \equiv \frac{\mathcal{E}_1(\zeta, \psi_1, \psi_2, \bar{\lambda})}{\mathcal{E}_0(\zeta, \psi_1, \psi_2, \bar{\lambda})}, \quad \mathcal{V}(\zeta, \psi_1, \psi_2, \bar{\lambda}) \equiv \frac{\mathcal{E}_2(\zeta, \psi_1, \psi_2, \bar{\lambda})}{\mathcal{E}_0(\zeta, \psi_1, \psi_2, \bar{\lambda})}.$$

Precise asymptotics

Theorem

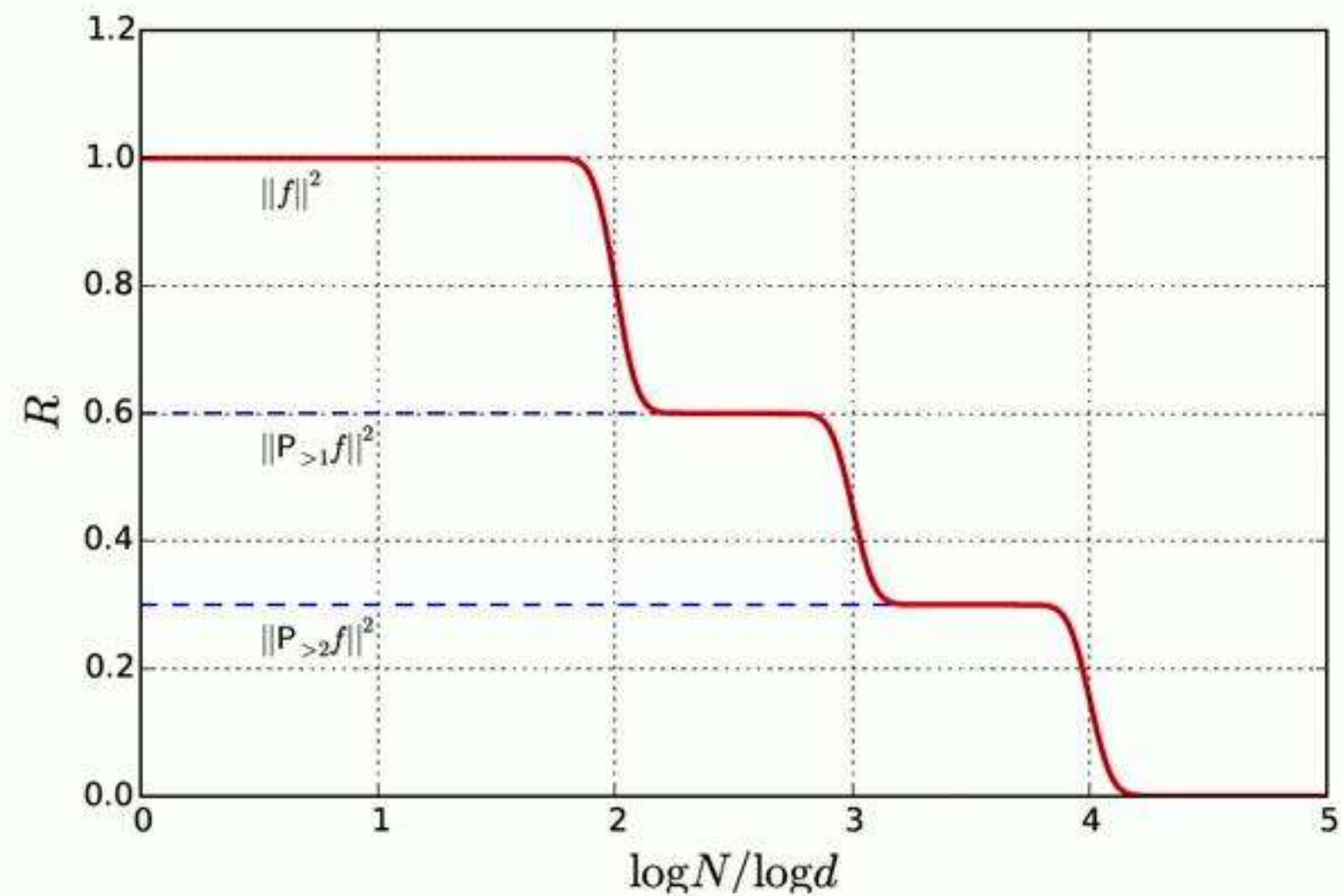
Assume $f_\star(x) = \langle \beta_0, x \rangle$ and define (for $G \sim \mathcal{N}(0, 1)$)

$$b_\star^2 = \mathbb{E}[\sigma(G)^2] - \mathbb{E}[\sigma(G)]^2 - b_1^2, \quad \zeta \equiv \frac{b_1^2}{b_\star^2}.$$

Then, for any $\lambda > 0$, we have

$$R_{\text{RF}}(\hat{f}_\lambda, f_\star) = \|\beta_0\|_2^2 \mathcal{B}(\zeta, \psi_1, \psi_2, \lambda/\bar{b}_\star^2) + \tau^2 \mathcal{V}(\zeta, \psi_1, \psi_2, \lambda/\bar{b}_\star^2) + o_d(1),$$

where $\mathcal{B}(\zeta, \psi_1, \psi_2, \bar{\lambda})$, $\mathcal{V}(\zeta, \psi_1, \psi_2, \bar{\lambda})$ are explicitly given below.



What happens at finite n ?