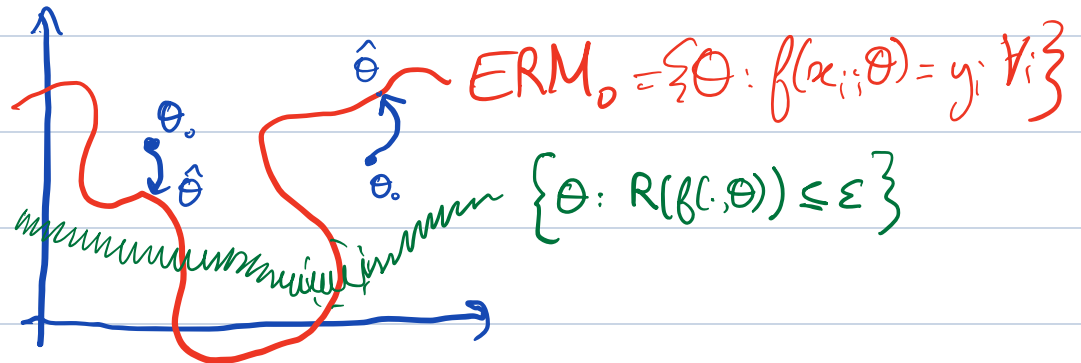


Lecture 4:Benign OverfittingLast week:

$\hat{\theta}$: depends on algo, initialization, parametrization etc
(step size etc.)

What about generalization of $f(\cdot; \hat{\theta})$?

→ model trained with no control on the model complexity
(until interpolation of the data)

≠ standard wisdom in statistical learning

→ models that are too complex will not generalize well

→ overfitting is bad and should be avoided

Today: "benign overfitting" → interpolating solution can generalize well

* Classical picture

* min- $\|\cdot\|_2$ interpolating solution in linear models (solution attained by GD)

↳ generalizes via "self-induced regularization"

* Example: inner-product kernel

→ RMT & exact asymptotics
→ matrix concentration

No time for these
→ will go back to it in lecture 6

Classical Picture

(2)

Parametric family: $\mathcal{F} = \{f(\cdot; \theta) : \theta \in \mathbb{R}^p\}$

Data: $\{(y_i, x_i)\}_{i=1}^m$ $x_i \in \mathbb{R}^d \stackrel{iid}{\sim} P$ $y_i \in \mathbb{R}$
 $y_i = f(x_i; \theta_*) + \varepsilon_i$, more e.g. $\varepsilon_i \stackrel{iid}{\sim} N(0, \sigma_\varepsilon^2)$

SRM: $\hat{\theta}(\lambda) = \underset{\theta \in \mathbb{R}^p}{\operatorname{argmin}} \left\{ \frac{1}{m} \sum_{i=1}^m (y_i - f(x_i; \theta))^2 + \lambda \|\theta\|_2^2 \right\}$

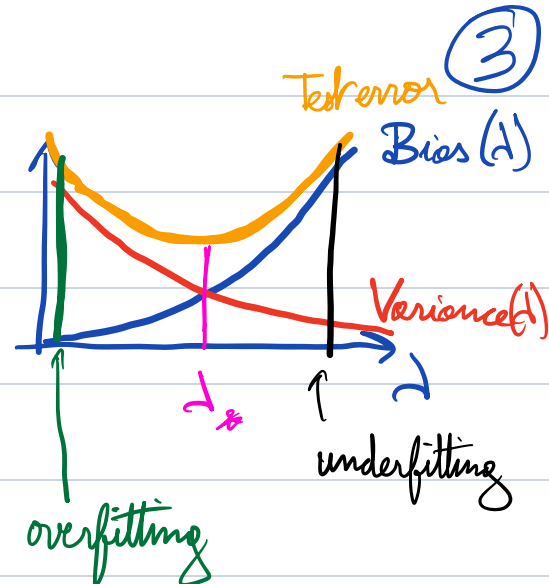
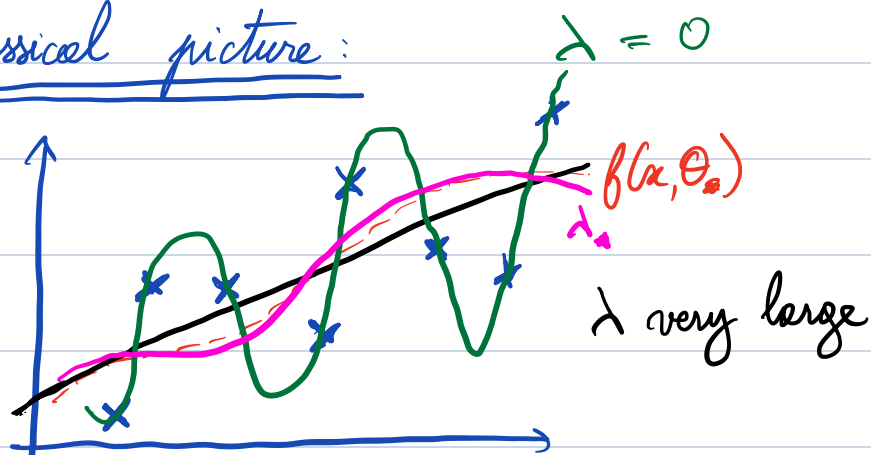
Errors

Test error: $R(\theta_*, \lambda) = \mathbb{E}_x \left[(f(x; \theta_*) - f(x; \hat{\theta}(\lambda)))^2 \right]$

Bias - Variance decomposition:

$$\begin{aligned} \mathbb{E}_\varepsilon [R(\theta_*, \lambda)] &= \mathbb{E}_{x, \varepsilon} \left[(f(x; \theta_*) - f(x; \hat{\theta}(\lambda)))^2 \right] \\ &= \mathbb{E}_x \left[(f(x; \theta_*) - \mathbb{E}_\varepsilon [f(x; \hat{\theta}(\lambda))])^2 \right] + \mathbb{E}_{x, \varepsilon} \left[(f(x; \hat{\theta}(\lambda)) - \mathbb{E}_\varepsilon [f(x; \hat{\theta}(\lambda))])^2 \right] \\ &=: \text{BIAS}(\lambda) + \text{VARIANCE}(\lambda) \\ &= \mathbb{E}_x \left[\text{bias}(\mathbb{E}_\varepsilon [f(x; \hat{\theta}(\lambda))])^2 \right] = \mathbb{E}_x \left[\text{Var}_\varepsilon (f(x; \hat{\theta}(\lambda))) \right] \end{aligned}$$

Classical picture:



Trade-off between Bias and Variance

→ need to carefully choose regularization parameter λ in order to control the complexity of the fitted model $f(x; \theta(\lambda))$

⇒ too complex: small bias, large variance

OVERFITTING

⇒ too simple: large bias, small variance

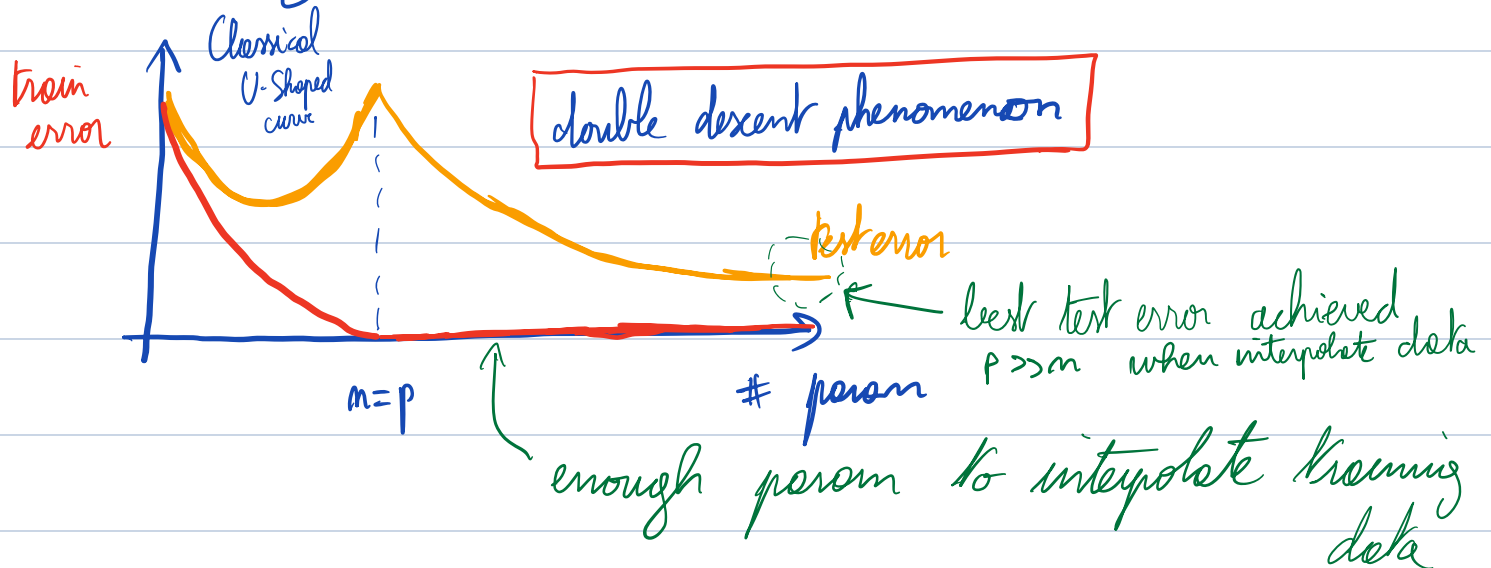
UNDERFITTING

Modern approach:

No careful control of the complexity of the model

(4)

In fact, generalizing well even when trained until interpolation of the training data $f(x_i; \hat{\theta}) = y_i = f(x_i; \theta_0) + \varepsilon_i$

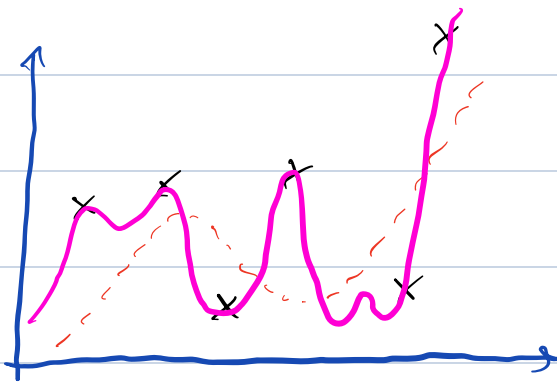


Today: I will present a scenario where the balance between the bias and the variance is achieved not by carefully tuning an explicit parameter but by a novel phenomena

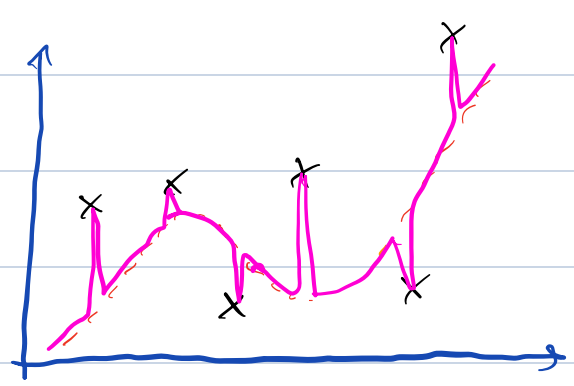
Self-induced regularization

⇒ Show sufficient & necessary conditions for benign overfitting in linear models

(5)



Expected picture
for overfitting



Sometimes
benign

Benign overfitting

only way to get $\hat{R}_n(\hat{\beta}) = 0$
 $R(\hat{\beta}) \ll 1$
when noise

estimator: $\hat{f}(x, \hat{\theta}) = \underbrace{\hat{f}_0(x, \hat{\theta}_0)}_{\downarrow} + \underbrace{\Delta(x)}_{\downarrow}$

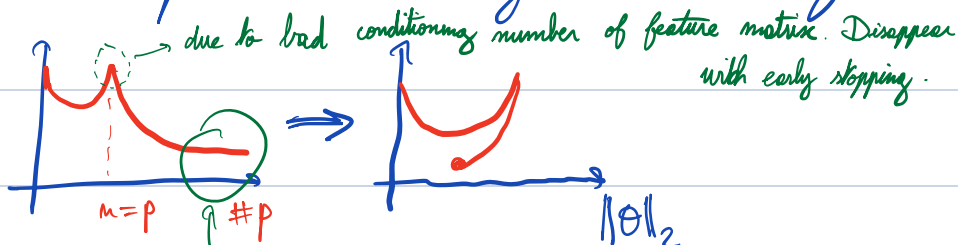
good for prediction
because smooth

spikes that are useful
for interpolation but do
not harm prediction
because $\mathbb{E}_x[\Delta(x)^2] \ll 1$.

(6)

Remark 1: Double descent: we saw that # param is not a good measure of model complexity

→ as $\|\theta\|_2$
recover classical U-shaped curve

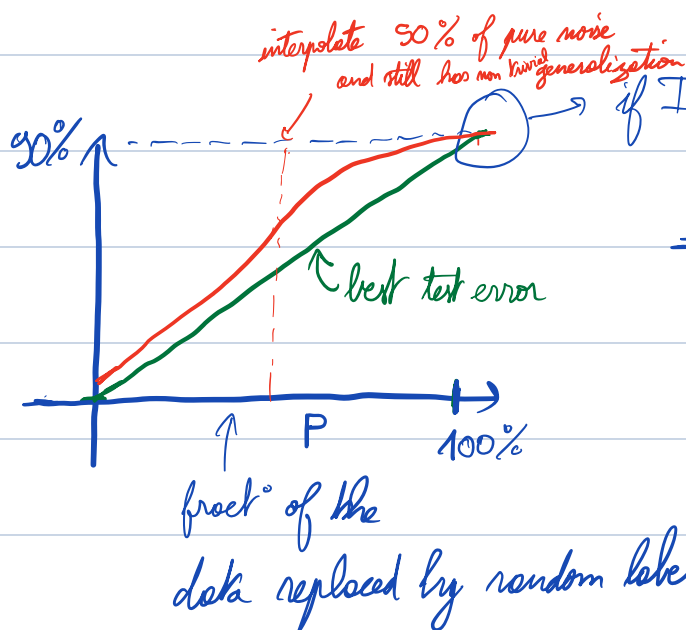


interesting phenomenon here is that we can interpolate noisy data and still generalize well.

Remark 2: Can UC explain benign overfitting?

→ bounds from lecture 2 independent of the data distribution

"Understanding deep learning requires rethinking generalization" Zhong et al. 2017



if I guess one of the 10 classes at random

⇒ same NN trained to interpolate data with different level of corrupted samples

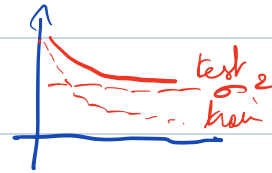
$R_{d_n}(F) = 1$ / can interpolate pure noise data

Simple rank: $y_i = f_*(x_i) + \varepsilon_i$

⑦

$$\hat{R}_n(\hat{f}) = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{f}(x_i))^2 \leq (1-\kappa) \sigma^2$$

$$R(\hat{f}) = \sigma^2 + \mathbb{E}[(\hat{f}(x) - f_*(x))^2]$$



$$\varepsilon_n(F) \geq |\hat{R}_n(\hat{f}) - R(\hat{f})| \geq \kappa \sigma^2$$

bounded away
from 0.

RIDGE REGRESSION

Setting: * Data $\{(y_i, x_i)\}_{i=1}^n$ $x_i \in \mathbb{R}^p$ $y_i \in \mathbb{R}$

* $\mathbb{E}[x_i] = 0$ $\mathbb{E}[x_i x_i^T] = \Sigma$

$$x_i = \Sigma^{\frac{1}{2}} z_i$$

* x_i are iid $y_i = \langle \theta_*, x_i \rangle + \varepsilon_i$

noise ε_i independent $\mathbb{E}[\varepsilon_i] = 0$ $\mathbb{E}[\varepsilon_i^2] = \sigma_\varepsilon^2$

Fit a linear model: $f(x, \theta) = \langle x, \theta \rangle$ $\theta \in \mathbb{R}^p$

Ridge regression: $\hat{\theta}(\lambda) = \underset{\theta \in \mathbb{R}^p}{\operatorname{argmin}} \left\{ \|y - X\theta\|_2^2 + \lambda \|\theta\|_2^2 \right\}$

$$X = \begin{matrix} \uparrow \quad \overbrace{\quad}^p \quad \downarrow \\ \begin{bmatrix} - & x_1 & - \\ & \vdots & \\ - & x_n & - \end{bmatrix} \end{matrix} \quad y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}$$

$$= X^T S y$$

where $S = (XX^T + \lambda \operatorname{Id})^{-1}$

Limit $\lambda \searrow 0^+$: $\hat{\theta}(0^+) = \underset{\theta \in \mathbb{R}^p}{\operatorname{argmin}} \left\{ \|\theta\|_2 : y = X\theta \right\}$

$$= X^T S y \quad S = (XX^T)^+$$

⇒ Minimum- $\|\cdot\|_2$ norm interpolating solution
 ↳ GD with $\theta^0 = 0$ converges to $\hat{\theta}(0^+)$

9

Excess test error: $R(\hat{\theta}) = \mathbb{E}_x[(f(x; \hat{\theta}) - f(x; \theta_*))^2]$
 $= \mathbb{E}_x[(\langle x, \hat{\theta} - \theta_* \rangle)^2]$
 $= \|\hat{\theta} - \theta_*\|_{\Sigma}^2$
 $= (\hat{\theta} - \theta_*)^T \Sigma (\hat{\theta} - \theta_*)$

$\|u\|_{\Sigma} = \|\Sigma^{\frac{1}{2}} u\|_2$

Bias-Variance decomposition:

$$\bar{R}(\hat{\theta}(\lambda)) = \mathbb{E}_{\varepsilon}[R(\hat{\theta}(\lambda))] = \mathbb{E}_{\varepsilon}[\|\hat{\theta}(\lambda) - \theta_*\|_{\Sigma}^2]$$

$$X^T S y = X^T S X \theta_* + X^T S \varepsilon$$

$$= B(\lambda) + V(\lambda)$$

where $B(\lambda) = \|(\hat{I}_d - X^T S X) \theta_*\|_{\Sigma}^2$

$$V(\lambda) = \sigma^2 \text{Tr}(S X \Sigma X^T S)$$

Self-induced regularization: high level explanation

WLOG: assume $\Sigma = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_p)$ $\lambda_1 \geq \lambda_2 \geq \dots$

Consider top k eigenspaces (corresponding to top k eigenvalues)
 $\lambda_1 \dots \lambda_k$

write $x = \begin{bmatrix} x_0 \\ x_+ \end{bmatrix}$ $\begin{matrix} \text{first } k \text{ coordinates} \\ \text{p-k last coordinates} \end{matrix}$

$$X = \begin{bmatrix} -x_1- \\ \vdots \\ -x_m- \end{bmatrix} = \begin{matrix} \uparrow \downarrow \\ m \end{matrix} \begin{matrix} \leftarrow k \rightarrow & \leftarrow p-k \rightarrow \\ X_0 & | & X_+ \end{matrix}$$

$$\Sigma = \begin{matrix} \uparrow k \downarrow \\ \left(\begin{array}{c|c} \Sigma_0 & \\ \hline & \Sigma_+ \end{array} \right) \\ \leftarrow p-k \rightarrow \end{matrix}$$

$$\Theta = (\Theta_0, \Theta_+) \\ \downarrow \qquad \qquad \downarrow \\ (\theta_1, \dots, \theta_k) \qquad (\theta_{k+1}, \dots, \theta_p)$$

$$XX^T = X_0 X_0^T + \underbrace{X_+ X_+^T}_{\text{assume } \approx \gamma \text{ Id}_m}$$

(later, assume $\frac{1}{c} \text{Id} \preceq M \preceq c \text{Id}$)

Recall: $\hat{\Theta}(\lambda) = X^T S y \approx X^T (X_0 X_0^T + (\gamma + \lambda) \text{Id})^{-1} y$

(11)

$$\rightarrow \hat{\Theta}_0 = X_0^\top (X_0 X_0^\top + (\lambda + \gamma) \text{Id})^{-1} y$$

top k coordinates

$$\rightarrow \text{the same as } \hat{\Theta}_0 = \underset{\Theta_0 \in \mathbb{R}^k}{\operatorname{argmin}} \left\{ \|y - X_0 \Theta_0\|_2^2 + (\lambda + \gamma) \|\Theta_0\|_2^2 \right\}$$

"self induced regularization" from the high-degree part of the features $X_+ X_+^\top$

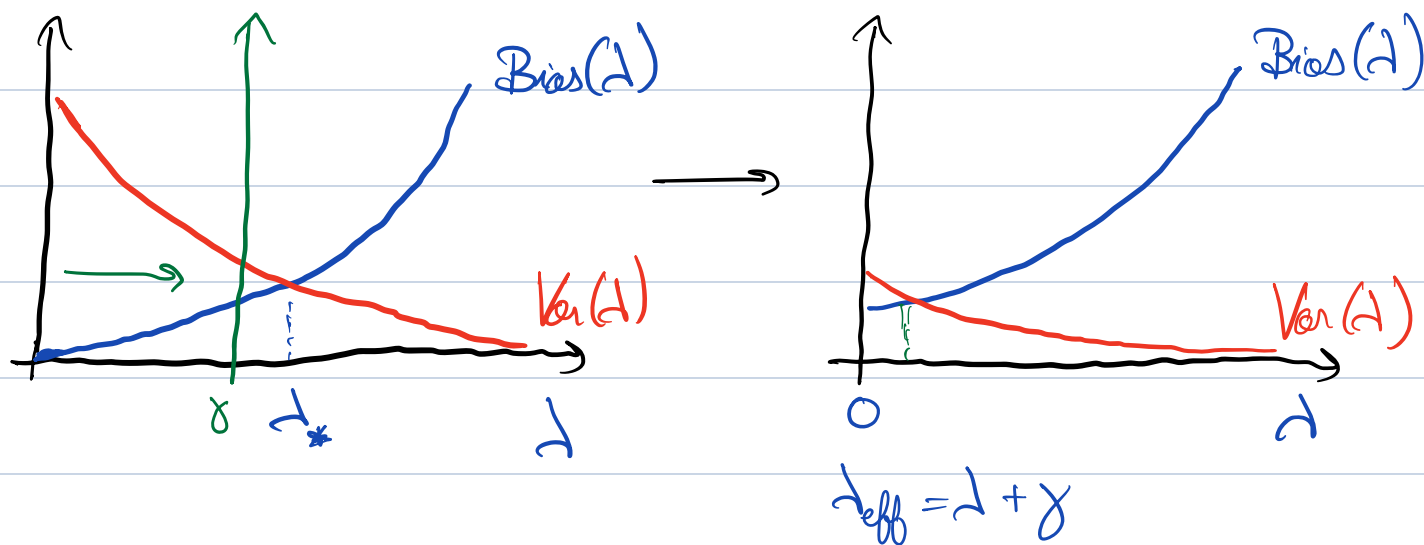
Even when $\lambda = 0^+$: $\hat{\Theta}_0$ solution of a lower-dim problem with effective regularization $\gamma > 0$.

* $\hat{\Theta}_+$ can show that it is small and does not contribute to the test error but interpolate the training data

$f(x, \hat{\Theta}) = \underbrace{\langle \hat{\Theta}_0, x_0 \rangle}_{\text{"signal part" of the features}} + \underbrace{\langle \hat{\Theta}_+, x_+ \rangle}_{\text{"noise part" of feature}}$
 $\nearrow \hat{f}_0(x)$ spiky part $\Delta(x)$
 → generalize well thanks to γ → interpolate data, doesn't contribute to test error

⇒ Noise part acts as an effective ridge regularization

Effective regularization $\lambda \rightarrow \lambda_{\text{eff}} = \lambda + \gamma$



"Shifted Bias-Variance curve"

$\lambda = 0$ close to optimal regularization

"Benign Overfitting in Ridge Regression"

- Bartlett, Trigler, 2020

13

$$x_i = \sum_{j=1}^n z_{ij} z_j \quad z_i \text{ iid sub-Gaussian} \quad (\text{for simplicity})$$

↓
more general

Recall: $S = (\lambda \text{Id} + XX^T)^{-1} = (\lambda \text{Id} + X_+ X_+^T + X_- X_-^T)^{-1}$

$S_+ = (\lambda \text{Id} + X_+ X_+^T)^{-1} \quad S = (S_+^{-1} + X_- X_-^T)^{-1}$

even for non linear kernels

TMM [Bartlett, Trigler, '20] Assume w.h.p. $\frac{\lambda_{\max}(S_+)}{\lambda_{\min}(S_+)} \leq L$

Then whp:

$$\text{Bias}(\lambda) \lesssim L^4 \left[\|\theta_0^*\|_{\Sigma_0^{-1}}^2 \left(\frac{\lambda + \sum_{i>k} \lambda_i}{n} \right)^2 + \|\theta_+^*\|_{\Sigma_+}^2 \right]$$

$$\text{Var}(\lambda) \lesssim \sigma_\varepsilon^2 L^2 \left[\frac{k}{n} + \frac{n \sum_{i>k} \lambda_i^2}{\left(\lambda + \sum_{i>k} \lambda_i \right)^2} \right]$$

RMK: * Matching lower bound (up to multiplicative constants / minimax)

* In particular, interpolators $\lambda=0$ can be near optimal.

* Does not fit θ_+^* at all

* Does not depend on p (if sums $\sum \lambda_i$ cv : applies to $p=\infty$)

B.O.: small variance: need to choose k so that

$$(1) \quad k \ll n \quad (2) \quad \frac{\left(\sum_{i>k} \lambda_i\right)^2}{\sum_{i>k} \lambda_i^2} \gg n$$

E.g. $d_{k+1} := \# \text{ eigenvalues of } \Sigma \text{ in } \left[\frac{\lambda_{k+1}}{2}, \lambda_{k+1} \right]$

and assume eigenvalues $< \frac{\lambda_{k+1}}{2}$ are negligible

$$\Rightarrow (2) \text{ requires } \frac{(d_{k+1} \cdot \lambda_{k+1})^2}{d_{k+1} \cdot \lambda_{k+1}^2} = \boxed{d_{k+1} \gg n}$$

and we have Variance $\lesssim \frac{k}{n} + \frac{n}{d_{k+1}}$

Variance of fitting Θ_0^*
(actually this is the parametric rate!)

variance of fitting the rest
effect of overfitting
 $\Downarrow$
negligible

Rmk. $\lambda_i = \frac{1}{i^\alpha (1 + \log i)^\beta}$ decay of eigenvalues

Test error ($\lambda = 0^+$) $\rightarrow 0$ as $n \rightarrow \infty$ "consistency of interpolating solution"

IFF $\alpha = 1$ $\beta > 1$

To get benign overfitting here:

(1) Overparametrization: $p \gg n$

directions in parameter space unimportant for prediction $\gg$ sample size

(2) Smallest eigenvalues of Σ must decay slowly

$d_{k+1} \gg n$

Self-induced regularization $\Sigma^{\frac{1}{2}} z_i$ z_i iid c -sub-G coordinates add way more generally

Lemma. $\exists c > 0$ s.t. with prob at least $1 - 2 \exp(-\frac{n}{c})$

$$\frac{\lambda_{\max}(X_+ X_+^T)}{\lambda_{\min}(X_+ X_+^T)} \leq c \frac{\sum_{i \geq k} d_i + n d_{k+1}}{\sum_{i \geq k} d_i - c n d_{k+1}} \leq L$$

$n d_{k+1} \ll \sum_{i \geq k} d_i$

Example: Kernel regression with inner product kernel on the sphere

$$u_i \underset{iid}{\sim} \text{Unif}(S^{d-1}(\sqrt{d}))$$

$$y_i = f_*(u_i) + \varepsilon_i$$

general squared integrable function

Linear model: $\hat{f}(u; \theta) = \langle \theta, \phi(u) \rangle$

$$\phi(u) = \begin{pmatrix} \zeta_0, \zeta_1 u_1, \dots, \zeta_1 u_d \\ \zeta_2 Y_{21}(u), \dots, \zeta_2 Y_{2d^2}(u) \\ \vdots \\ \zeta_k Y_{k1}(u), \dots, \zeta_k Y_{kO(d^k)}(u) \\ \vdots \end{pmatrix}$$

basis of degree k polynomials in u .
 $O(d^k)$

$\zeta_k = O(d^{-\frac{k}{2}})$

$$x_i = \phi(u_i) = \Sigma^{\frac{1}{2}} z_i \quad \mathbb{E}[z_i z_i^T] = \text{Id}_{\infty}$$

$$\Sigma = \begin{pmatrix} \zeta_0^2 & & & \\ & \zeta_1^2 & & \\ & & \zeta_1^2 & \\ & & & \zeta_2^2 & \\ & & & & \zeta_2^2 & \\ & & & & & \ddots \end{pmatrix}$$

$\xleftarrow{d} \quad \xleftarrow{O(d^2)}$
 $\searrow \quad \searrow$

$$\text{In } d^l \ll n \ll d^{l+1}$$

Take $k := d^l$ $X_+ X_+^T \approx \left(\sum_{\Delta=l+1}^{\infty} B_{d,\Delta} \Xi_{\Delta}^2 \right) \text{Id}_n$

$$f_*(u) = \underbrace{\langle \alpha_0, \Theta_{*,0} \rangle}_{\text{best degree-}l \text{ polynomial approx}} + \langle \alpha_+, \Theta_{*,+} \rangle$$

$$\hat{\Theta}_0 \approx \Theta_{*,0} \quad \mathbb{E}_{\alpha} [\langle \hat{\Theta}_+, \alpha_+ \rangle^2] = \|\hat{\Theta}_+\|_{\Sigma_+}^2 \approx 0$$

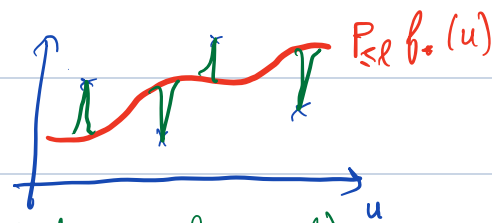
[needs very different tools to analyze $\rightarrow$ not sub-Gaussian]

Thm [M., Ghorbani, Mei, Montanari, '19]
 If $d^{l+\delta} \leq n \leq d^{l+1-\delta}$ as $n, d \rightarrow \infty$, then

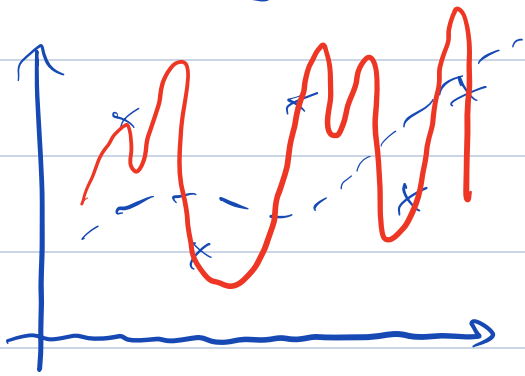
$$\text{Test error} = \|P_{\geq l} f_*\|_{L^2}^2 + o(1)$$

$$\hat{f}(u) = P_{\leq l} f_*(u) + \Delta(u)$$

spiky part (high degree polynomial)



Conclusion:

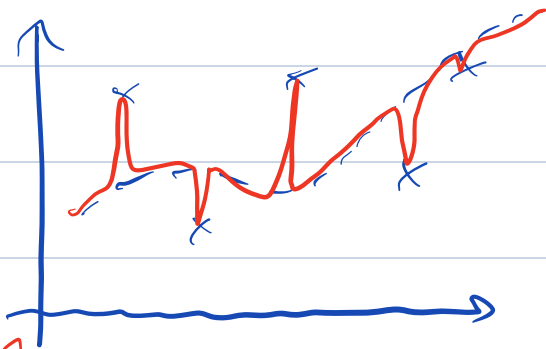


classical picture

BAD OVERFITTING



Modern picture



sometimes

BENIGN OVERFITTING

$$\hat{b} = b_0 + \Delta$$

Above "low-dim" ^{classical} picture misleading

→ much easier to interpolate benignly
in high dim

(Of course, it required a lot of work to realize it)

Proof of Bartlett / Trigler

$$S = (X_0 X_0^T + X_+ X_+^T + I)^{-1}$$

$$S_+ = (X_+ X_+^T + I)^{-1}$$

$$\frac{\lambda_{\max}(S_+)}{\lambda_{\min}(S_+)} \leq L$$

Here only prove Variance term ($\sigma_\varepsilon^2 = 1$ WLOG)

$$V(X, I) = \text{Tr}(S X \Sigma X S) = V_0 + V_+$$

$$V_0 := \text{Tr}(S X_0 \Sigma_0 X_0^T S)$$

$$V_+ := \text{Tr}(S X_+ \Sigma_+ X_+^T S)$$

1st Term:

$$V_0 = \text{variance along top } k \text{ directions}$$

$$\lesssim L^2 \frac{k}{n}$$

Proof:

$$S X_0 = (X_0 X_0^T + S_+^{-1})^{-1} X_0$$

$$= S_+ X_0 (I + X_0^T S_+ X_0)^{-1}$$

Introduce $M := \Sigma_0^{\frac{1}{2}} (I + X_0 S_+ X_0^T)^{-1} \Sigma_0^{\frac{1}{2}}$

$$V_0 = \text{Tr}(S_+ \overbrace{X_0 \Sigma_0^{-\frac{1}{2}}}^{= Z_0} M^2 \Sigma_0^{-\frac{1}{2}} X_0^T S_+)$$

$$V_0 \leq \|M\|_{op}^2 \operatorname{Tr}(S_+ Z_0 Z_0^T S_+)$$

$$\begin{aligned} \|M\|_{op} &= \|(\Sigma_0^{-1} + \Sigma_0^{-\frac{1}{2}} X_0^T S_+ X_0 \Sigma_0^{-\frac{1}{2}})\|_{op} \leq \frac{1}{\lambda_{\min}(Z_0^T S_+ Z_0)} \\ &\leq \frac{1}{\lambda_{\min}(S_+) \lambda_{\min}(Z_0^T Z_0)} \end{aligned}$$

Mence

$$V_0 \leq \frac{\lambda_{\max}(S_+)^2 \operatorname{Tr}(Z_0^T Z_0)}{\lambda_{\min}(S_+)^2 \cdot \lambda_{\min}(Z_0^T Z_0)^2}$$

$$\frac{\operatorname{Tr}(A)}{\lambda_{\min}(A)} \leq \frac{\lambda_{\max}(A)}{\lambda_{\min}(A)} \cdot \operatorname{rank}(A)$$

$$\leq L^2 \frac{\lambda_{\max}(Z_0^T Z_0)}{\lambda_{\min}(Z_0^T Z_0)} \times \frac{k}{\lambda_{\min}(Z_0^T Z_0)}$$

$$\left[\begin{array}{l} Z_0 \in \mathbb{R}^{m \times p} \text{ iid } O(1)\text{-subGaussian} \quad k \leq c n \\ \lambda_{\min}(Z_0), \lambda_{\max}(Z_0) \asymp \sqrt{n} \end{array} \right]$$

$$V_0 \lesssim L^2 \cdot \frac{k}{n}$$

□

2nd term $V_+ = \text{variance due to overfitting}$

$$\leq L^2 \frac{\sum_{i>k}^m d_i^2}{\left(\lambda + \sum_{i>k} d_i\right)^2}$$

Rank: $\frac{\lambda_{\max}(S_+)}{\lambda_{\min}(S_+)} \leq L \Rightarrow \lambda_{\max}(S_+) \leq \frac{L}{\lambda + \sum_{i>k} d_i}$

$$\lambda_{\max}(S_+) \leq L \lambda_{\min}(S_+) \leq L \frac{n}{\text{Tr}(S_+^{-1})}$$

$$\frac{1}{n} \text{Tr}(S_+^{-1}) = \frac{1}{n} \text{Tr}(\lambda \text{Id} + X_+ X_+^T) = \lambda + \frac{1}{n} \sum_{i=1}^n \|\alpha_{i,+}\|_2^2$$

WLLN

$$\asymp \lambda + \text{Tr}(\Sigma_+)$$

Proof: $V_+ = \text{Tr}(S X_+ \Sigma_+ X_+^T S) \leq \lambda_{\max}(S)^2 \text{Tr}(X_+ \Sigma_+ X_+^T)$

$$\leq \lambda_{\max}(S_+)^2 \sum_{i=1}^n \langle z_{i,+}, \Sigma_+^2 z_{i,+} \rangle$$

$$\leq \lambda_{\max}(S_+)^2 n \cdot \text{Tr}(\Sigma_+^2)$$

Using the above remark

$$V_+ \leq L^2 \frac{n \cdot \text{Tr}(\Sigma_+^2)}{(\lambda + \sum_{i \geq k} \lambda_i)^2} \quad \square$$

Sharp asymptotics using RMT

So far we only used matrix concentration

↳ in particular: bounds hold for $n \geq C$.

⇒ Lots of work have considered to derive sharp formulas for the bias and variance as $n, p \rightarrow \infty$ using RMT

Below a very brief presentation of these results

[Mastie, Montanari, Rosset, Tibshirani, 2019]

(Follow Andrea's presentation)

Setting: $\{(\alpha_i, y_i)\}_{i=1}^n$ $y_i \in \mathbb{R}$ $\alpha_i \in \mathbb{R}^p$ iid

$$* y_i = \langle \theta_*, \alpha_i \rangle + \varepsilon_i \quad \mathbb{E} \varepsilon_i = 0 \quad \mathbb{E} \varepsilon_i^2 = \sigma_\varepsilon^2$$

$$* \alpha_i = \Sigma^{\frac{1}{2}} z_i \quad \mathbb{E} z_i = 0 \quad \mathbb{E}[z_i z_i^T] = I_p$$

$$* z_i \text{ independent coordinates } \mathbb{E}[|z_{ij}|^k] \leq C_k < \infty \quad \forall i, k$$

$$\text{Ridge solution } \hat{\theta}(\lambda) = \underset{\theta \in \mathbb{R}^p}{\operatorname{argmin}} \left\{ \frac{1}{n} \|y - X\theta\|_2^2 + \lambda \|\theta\|_2^2 \right\} = \frac{1}{n} X^T (\lambda I + \frac{X X^T}{n})^{-1} y$$

$$R(\hat{\theta}) = \|\hat{\theta}(\lambda) - \theta_*\|_2^2 = \mathcal{B}_n + V_n$$

$$B_m = \lambda^2 \langle \theta_*, S_\lambda \Sigma S_\lambda \theta_* \rangle \quad \hat{\Sigma} = \frac{1}{n} X^T X$$

$$V_m = \frac{\sigma^2}{n} \text{Tr}(\Sigma \hat{\Sigma} S_\lambda^2) \quad S_\lambda = (\hat{\Sigma} + \lambda I)^{-1}$$

Then [Hastie et al, 19] $\Sigma = I \quad \lambda = 0^+$ (min-norm)

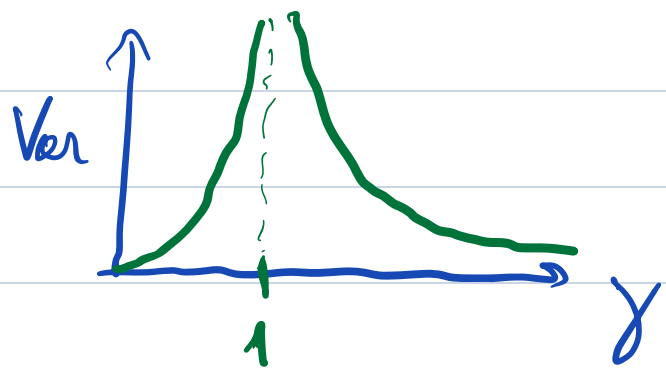
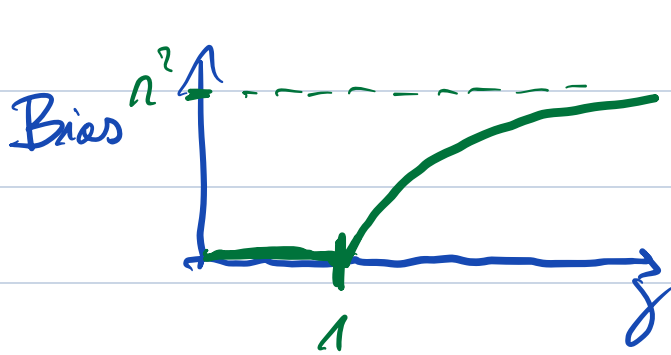
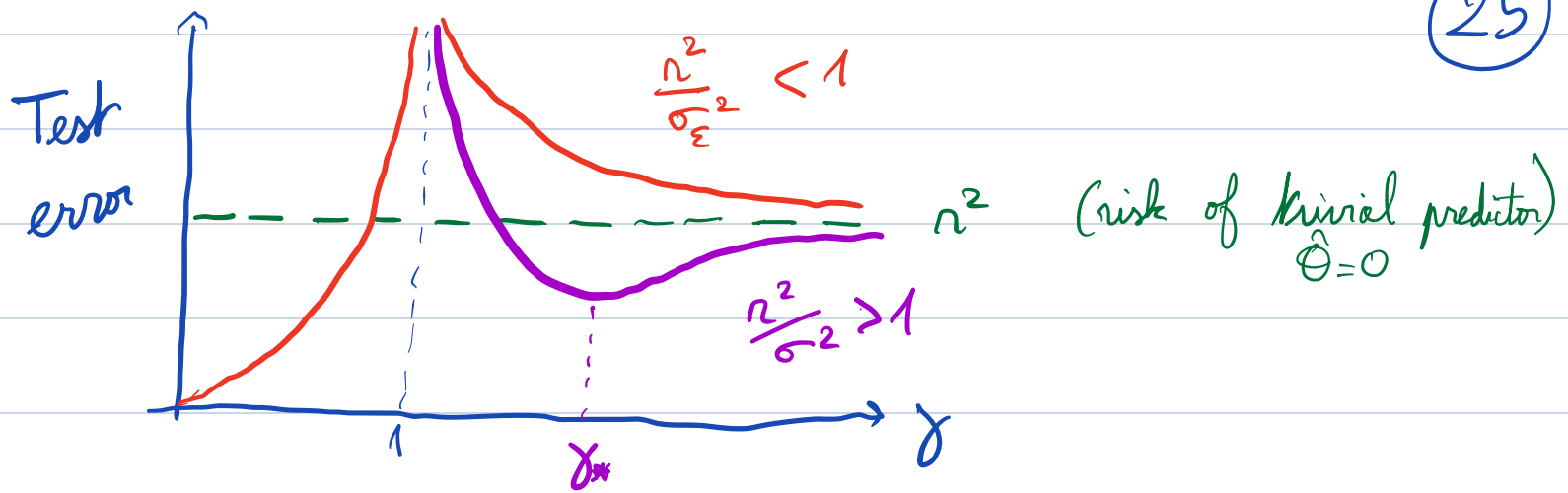
$$n, p \rightarrow \infty \quad \frac{p}{n} \rightarrow \gamma \in (0, \infty) \quad \|\theta_*\|_2^2 \rightarrow n^2$$

$$\text{Then } \lim_{n, p \rightarrow \infty} B_m = n^2 \left(1 - \frac{1}{\gamma}\right)_+$$

$$\lim_{n, p \rightarrow \infty} V_m = \sigma_\varepsilon^2 \frac{1}{\left(\frac{1}{\gamma} \vee \gamma - 1\right)}$$

In particular,

$$\lim_{n, p \rightarrow \infty} R(\hat{\theta}) = \begin{cases} \frac{\gamma \sigma_\varepsilon^2}{1 - \gamma} & \text{if } \gamma < 1 \\ n^2 \left(1 - \frac{1}{\gamma}\right) + \frac{\sigma^2}{\gamma^{-1}} & \text{if } \gamma > 1 \end{cases}$$



- Remark:
- * Prove way more stuff in their paper
 - * Double descent
 - however not minimizing $\gamma \rightarrow \infty$
 - * Not a good model $p = d$
 - ↳ true fct $f_*(x) = \langle \Theta_*, x \rangle$ becomes more complex as p grows
 - * Better models
 - [Mei, Montanari, '19]

Crash course Random Matrix Theory (RMT)

$$\hat{\Sigma} = \frac{XX^T}{n} \in \mathbb{R}^{n \times n} \quad \text{empirical spectral distribution (ESD)}$$

$$\mu_n = \frac{1}{p} \sum_{i=1}^p \delta_{\lambda_i(\hat{\Sigma})}$$

Marchenko - Pastur law: $\mu_n \xrightarrow[n/p \rightarrow \infty]{\text{(weak cv)}} \mu_\infty$
 $n/p \rightarrow \gamma$

To prove it: Stieltjes transform $z \in \mathbb{C}$

$$M_n(z) = \frac{1}{p} \sum_{i=1}^p \frac{1}{\lambda_i(\hat{\Sigma}) - z} = \frac{1}{p} \text{tr}((\hat{\Sigma} - zI_p)^{-1})$$

Thm: [MP] Let $z \mapsto m(z)$ be the analytic fct on $z \in \mathbb{C}_+$ s.t. $|m(z)| \leq \frac{C}{|z|}$ as $|z| \geq C$ and solution of

$$\frac{m}{1+z m} = 1 + \gamma m$$

Then a.s. $M_n(z) \rightarrow m(z)$ as $n, p \rightarrow \infty$
 $n/p \rightarrow \gamma$

Stieltjes inversion:

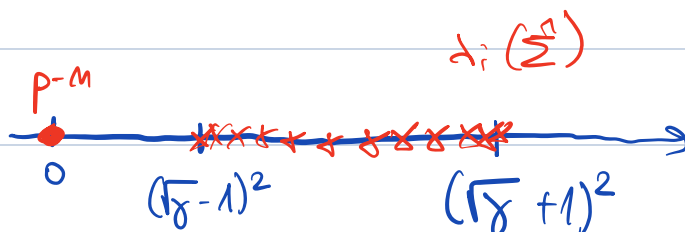
$$\mu([a, b]) = \lim_{\eta \rightarrow 0} \int_a^b \frac{1}{\pi} \operatorname{Im}(m(\lambda + i\eta)) d\lambda$$

Remark: $m(z) = \frac{1 - z - \gamma - \sqrt{(1 - z - \gamma)^2 - 4\gamma z}}{2\gamma z}$ by solving the fixed pt equation

Thm: [Bei-Yin law] $\gamma > 1$ $\operatorname{rank}(\hat{\Sigma}) = n$ w.h.p.

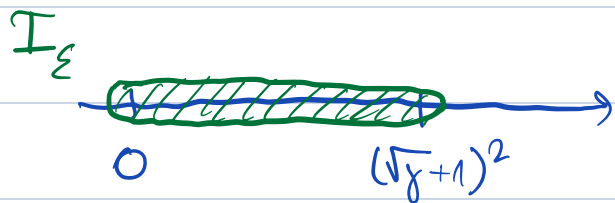
$$\begin{cases} \lambda_{\max}(\hat{\Sigma}) = (\sqrt{\gamma} + 1)^2 + o_p(1) \\ \lambda_{\min}(\hat{\Sigma}) = (\sqrt{\gamma} - 1)^2 + o_p(1) \end{cases}$$

For $\gamma \leq 1$ $\lambda(\hat{\Sigma}) = 0$ $p - n$



Remark: $M_n(z)$ is unif. cont. (differentiable)
on $\mathbb{C} \setminus I_\varepsilon$

$$I_\varepsilon = \left\{ z: \begin{array}{l} \operatorname{Re}(z) \in [-\varepsilon, (\sqrt{8}+1)^2 + \varepsilon] \\ |\operatorname{Im}(z)| \leq \varepsilon \end{array} \right\}$$



All derivatives unif. bounded
outside I_ε

(convince yourself using expression of $M_n(z)$)

$$M_n(z) \xrightarrow{\text{a.s.}} m(z) \quad \forall z \notin I_\varepsilon.$$

and same for derivatives

$$\text{Take } z = -\lambda \quad M_n(-\lambda) = \frac{1}{p} \operatorname{Tr}((\hat{\Sigma} + \lambda)^{-1}) \rightarrow m(-\lambda)$$

Proof of Thm:

$$\begin{aligned}
\boxed{\text{Variance}} \quad \frac{1}{\sigma_\varepsilon^2} V_m &= \frac{1}{n} \text{tr}(\hat{\Sigma}(\hat{\Sigma} + \lambda)^{-2}) \\
&= \frac{p}{n} \left[\frac{1}{p} \text{tr}((\hat{\Sigma} + \lambda)^{-1}) - \frac{1}{p} \lambda \text{tr}((\hat{\Sigma} + \lambda I)^{-2}) \right] \\
&= \frac{p}{n} \left[M_m(-\lambda) + \lambda \frac{\partial}{\partial \lambda} M_m(-\lambda) \right] \\
&\rightarrow \gamma \left[m(-\lambda) + \lambda \frac{\partial}{\partial \lambda} m(-\lambda) \right] \\
&= \gamma \frac{\partial}{\partial \lambda} (\lambda m(-\lambda))
\end{aligned}$$

$$m(-\lambda) = \frac{-1 - \lambda - \gamma + \sqrt{(1 + \lambda - \gamma)^2 + 4\lambda\gamma}}{2\gamma\lambda}$$

Ridgeless limit: $\lambda \downarrow 0$ $\gamma > 1$ $m(-\lambda) = \frac{\gamma-1}{\gamma\lambda} + \frac{1}{\gamma(\gamma-1)} + O(\lambda)$

$$\frac{\partial}{\partial \lambda} (\lambda m(-\lambda)) = \frac{1}{\gamma(\gamma-1)} + O(\lambda)$$

$$\lim_{\lambda \rightarrow 0} \lim_{n, p \rightarrow \infty} V_m(\lambda) = \frac{\sigma_\varepsilon^2}{\gamma-1} \quad \gamma > 1$$

Remark: * need to invert the limits (doable with some work)
 * same for $\gamma < 1$

Bias

$$B_m = \lambda^2 \langle \theta_*, (\hat{\Sigma} + \lambda)^{-2} \theta_* \rangle$$

κ isotropic: $B_m = \lambda^2 \|\theta_*\|_2^2 \frac{1}{p} \text{Tr}((\hat{\Sigma} + \lambda)^{-2}) + o_p(1)$

[e.g. prove it for $\kappa \sim N(0, \text{Id}_d)$]

$$B_m = \lambda^2 \|\theta_*\|_2^2 \left(-\frac{\partial}{\partial \lambda} M_m(-\lambda) \right)$$

$$\rightarrow \lambda^2 \left(1 - \frac{1}{\lambda} \right)_+ \quad \text{with some calculus}$$



Anisotropic case

$\Sigma \neq I$ risk will depend on both Σ, θ_*

Interpolating solution $\lambda = 0^+$

"Effective regularization" λ_* unique positive solution of

$$\frac{1}{\gamma} = \frac{1}{\rho} \operatorname{Tr}(\Sigma(\Sigma + \lambda_* I)^{-1})$$

$$\overline{B}_n := \frac{\lambda_*^2 \langle \theta_*, (\Sigma + \lambda_*)^{-2} \Sigma \theta_* \rangle}{1 - \frac{1}{n} \operatorname{Tr}(\Sigma^2 (\Sigma + \lambda_*)^{-2})}$$

$$\overline{V}_n := \frac{\sigma^2}{n} \frac{\operatorname{Tr}(\Sigma^2 (\Sigma + \lambda_*)^{-2})}{1 - \frac{1}{n} \operatorname{Tr}(\Sigma^2 (\Sigma + \lambda_*)^{-2})}$$

Thm [Mertik et al, '19] $\left[\begin{array}{l} \lim_{n, p \rightarrow \infty} |B_n - \overline{B}_n| = 0 \\ \lim_{n, p \rightarrow \infty} |V_n - \overline{V}_n| = 0 \end{array} \right. \quad \text{a.s.}$

Interpretation: test error:

(32)

Ridge regression $\longleftrightarrow$ gaussian sequence model

$$y^s = \Sigma^{\frac{1}{2}} \theta_* + \frac{\omega}{\sqrt{n}} g \quad g \sim N(0, I_p)$$

$$\hat{\theta}^s = \underset{\theta}{\operatorname{argmin}} \left\{ \|y^s - \Sigma^{\frac{1}{2}} \theta\|_2^2 + \lambda_* \|\theta\|_2^2 \right\}$$

ω fixed pt of: $\omega^2 = \frac{\sigma_\varepsilon^2}{n} + \mathbb{E}_g \left[\|\hat{\theta}^s - \theta_*\|_\Sigma^2 \right]$

$$R_n^s = \mathbb{E} \left[\|\hat{\theta}^s - \theta_*\|_\Sigma^2 \right] = \underbrace{B_n^s}_{\overline{B}_n} + \underbrace{V_n^s}_{\overline{V}_n}$$

	Original problem	Sequence model
design	X random	$\Sigma^{\frac{1}{2}}$ deterministic
reg.	$\lambda = 0^+$	$\lambda_* > 0 !!$
noise	σ_ε^2	ω^2

self induced regularization!! more regularization in the model than what is naively expected from $\lambda=0$.

Appendix:

Concentration of $X_+ X_+^T$

33

I'll present a nice decoupling argument that applies under quite general conditions
(much beyond sub-Gaussian e.g. inner-product kernels)

$$\{u_i\}_{i \in [m]} \text{ iid} \quad K = (K(u_i, u_j))_{i, j \in [m]} \in \mathbb{R}^{m \times m}$$

$K: \mathcal{U} \times \mathcal{U} \rightarrow \mathbb{R}$ p.s.d. fct $\rightarrow$ p.s.d operator $|K|$
(lecture on kernels)

$$K(u_i, u_j) = \sum_{k=1}^{\infty} \underbrace{\lambda_k^2}_{\text{eigenvalues}} \underbrace{\phi_k(u_i) \phi_k(u_j)}_{\text{eigenvectors}}$$

EIGENDECOMPOSITION

$$\lambda_1^2 \geq \lambda_2^2 \geq \dots \quad \mathbb{E}[\phi_k(u) \phi_\ell(u)] = \delta_{k\ell}$$

$$\|K\|_{\text{op}} = \lambda_1^2$$

$$\begin{aligned} \text{Tr}(|K|) &= \sum_{k=1}^{\infty} \lambda_k^2 < \infty \quad (\text{"trace-class"}) \\ &= \mathbb{E}[K(u, u)] \end{aligned}$$

Assumptions:

(A) Diagonal: $\mathbb{E} \left[\max_{i \in [n]} |K(u_i, u_i) - \text{Tr}(K)| \right] \leq C \sqrt{n \|K\|_{\text{op}} \text{Tr}(K)}$

(B) Hypercontractivity: $\forall p$ integer ≥ 2 , $\exists C_p$ s.t.

$$\forall v \in \mathbb{R}^n \quad \mathbb{E} \left[\left(\sum_k v_k \phi_k(u) \right)^p \right]^{\frac{1}{p}} \leq C_p \|v\|_2$$

Thm: [M., Mei, Montanari, 2021] For all $\delta > 0$, $\exists C$ s.t.

$$\mathbb{E} [\|K - \text{Tr}(K)I_n\|_q] \leq C \sqrt{n^{1+\delta} \|K\|_{\text{op}} \text{Tr}(K)}$$

Rmk: * Bounds on expectation: use Markov to get w.h.p

↳ to get better dependency in proba $\rightarrow$ subG: union bound

can also improve $n^{1+\delta} \rightarrow n$

$\rightarrow$ hyper: moment method

* $(\phi_k(u_i))_{k \geq 1} \rightarrow z_i$ iid sub-Gaussian entries

$\rightarrow$ easy to check (A) & (B)

$$\begin{aligned} * \quad \lambda_{\max}(K) &\leq \text{Tr}(K) \left(1 + C \sqrt{\frac{n^{1+\delta} \|K\|_{\text{op}}}{\text{Tr}(K)}} \right) \\ \lambda_{\min}(K) &\geq \text{Tr}(K) \left(1 - C \sqrt{\frac{n^{1+\delta} \|K\|_{\text{op}}}{\text{Tr}(K)}} \right) \end{aligned}$$

We will use the following result in random matrix concentration

Prop [Vershynin 2010: Theorem 5.48] Let $A \in \mathbb{R}^{n \times k}$ with

$A = [a_1, \dots, a_n]^T$ where $a_i \in \mathbb{R}^k$ are indep random vectors with common second moment matrix $\Sigma = \mathbb{E}[a_i a_i^T]$.

Let $\Gamma = \mathbb{E}[\max_{i \in [n]} \|a_i\|_2^2]$. Then there exists a universal constant C such that

$$\mathbb{E}[\|A\|_{\text{op}}^2]^{\frac{1}{2}} \leq (\|\Sigma\|_{\text{op}} \cdot n)^{\frac{1}{2}} + C \cdot (\Gamma \cdot \log(n \wedge k))^{\frac{1}{2}}$$

Proof of Theorem: $K = \text{ddiag}(K) + \Delta$

where $\text{ddiag}(K) = \text{diag}(K(u_i, u_i))_{i \in [n]}$

$$\Delta_{ij} = \begin{cases} K(u_i, u_j) & i \neq j \\ 0 & \text{if } i = j. \end{cases}$$

so that $\mathbb{E}[\|K - \text{Tr}(K) \cdot I_n\|_{\text{op}}]$

$$\leq \underbrace{\mathbb{E}[\|\text{ddiag}(K) - \text{Tr}(K) \cdot I_n\|_{\text{op}}]}_{(I)} + \underbrace{\mathbb{E}[\|\Delta\|_{\text{op}}]}_{(II)}$$

Using assumption (A):

$$(I) = \mathbb{E} \left[\max_{i \in [n]} |K(u_i, u_i) - \text{Tr}(K)| \right] \leq \sqrt{n \|K\|_{\text{op}} \text{Tr}(K)}$$

For the second term: note that Δ does not have independent rows or columns and we can't use the above proposition directly.

We will use a standard decoupling argument:

Lemma: Let $\Delta_{T_1, T_2} = (\Delta_{ij})_{i \in T_1, j \in T_2} \in \mathbb{R}^{|T_1| \times |T_2|}$

$$\mathbb{E} [\|\Delta\|_{\text{op}}] \leq 4 \max_{T \subseteq [n]} \mathbb{E} [\|\Delta_{T, T^c}\|_{\text{op}}]$$

Proof of lemma: For $\|v\|_2 = 1$

$$v^T \Delta v = \sum_{i \neq j} v_i \Delta_{ij} v_j = 4 \mathbb{E}_T \left[\sum_{\substack{i \in T \\ j \in T^c}} v_i \Delta_{ij} v_j \right]$$

where T is a random subset of $\{1, \dots, n\}$ where each element is selected with prob $\frac{1}{2}$

Therefore $\mathbb{E}[\|\Delta\|_{\text{op}}] = \mathbb{E}\left[\sup_{v \in S^m} v^T \Delta v\right]$

$$\leq 4 \mathbb{E}_T \mathbb{E}\left[\sup_{v \in S^m} \sum_{\substack{i \in T \\ j \in T^c}} v_i \Delta_{ij} v_j\right] \leq 4 \sup_{T \subseteq [n]} \mathbb{E}[\|\Delta_{T, T^c}\|_{\text{op}}]$$

□

Using this lemma, it suffices to bound $\|\Delta_{T, T^c}\|_{\text{op}}$ for all $T \subseteq [n]$

$$\Delta_{T, T^c} = (\Delta_{ij})_{i \in T, j \in T^c}$$

Each row is $\Delta_{i, T^c} = (K(u_i, u_j))_{j \in T^c}$

so conditional on $(u_j)_{j \in T^c}$, Δ_{T, T^c} has iid rows!!

Denote $\mathbb{E}_T$ the expectation over $(u_i)_{i \in T}$ conditional on $(u_j)_{j \in T^c}$. We can use the above general concentration bound

$$\mathbb{E}_T[\|\Delta_{T, T^c}\|_{\text{op}}] \leq \{\|\Sigma_T\|_{\text{op}} |T|\}^{\frac{1}{2}} + C \{\Gamma_T \cdot \log(|T| \wedge |T^c|)\}^{\frac{1}{2}}$$

where we denoted

$$\Sigma_T := \mathbb{E}_{u_i}[\Delta_{i, T^c} \Delta_{i, T^c}^T]$$

$$\Gamma_T := \mathbb{E}_T\left[\max_{i \in [n]} \|\Delta_{i, T^c}\|_2^2\right]$$

Mence

$$\mathbb{E}[\|\Delta\|_{op}] \leq C \sup_{T \subseteq [n]} \left\{ \left(\mathbb{E}_{T^c}[\|\Sigma_T\|_{op}] \cdot n \right)^{\frac{1}{2}} + C \left(\mathbb{E}_{T^c}[\Gamma_T] \cdot \log n \right)^{\frac{1}{2}} \right\}$$

Bounding $\mathbb{E}_{T^c}[\|\Sigma_T\|_{op}]$

$$\|\Sigma_T\|_{op} = \left\| \mathbb{E}_{u_i} [\Delta_{i,T^c} \Delta_{i,T^c}^T] \right\|_{op} = \sup_{\|v\|_2=1} \sum_{i,j \in T^c} \sum_k \xi_k^4 \phi_k(u_i) \phi_k(u_j) v_i v_j$$

$$\leq \|K\|_{op} \sup_{\|v\|_2=1} \sum_{i,j \in T^c} \xi_k^2 \phi_k(u_i) \phi_k(u_j) v_i v_j$$

$$\leq \|K\|_{op} \left\| (K(u_i, u_j))_{i,j \in T^c} \right\|_{op}$$

$$\leq \|K\|_{op} (\| \text{ddiag } K \|_{op} + \|\Delta\|_{op})$$

By hypercontractivity:

$$\mathbb{E}[\| \text{ddiag } K \|_{op}] \leq \mathbb{E} \left[\max_{i \in [n]} K(u_i, u_i)^p \right]^{\frac{1}{p}} \leq \mathbb{E} \left[\sum_{i=1}^n K(u_i, u_i)^p \right]^{\frac{1}{p}}$$

$$\leq n^{\frac{1}{p}} \mathbb{E} \left[K(u_i, u_i)^p \right]^{\frac{1}{p}} \leq C_p n^{\frac{1}{p}} \text{Tr}(K)$$

Hence

$$\mathbb{E}_{T^c}[\Sigma_T] \leq C_p n^{\frac{1}{p}} \|K\|_{op} \mathcal{H}(K) + \|K\|_{op} \mathbb{E}[\|\Delta\|_{op}]$$

Bounding $\mathbb{E}_{T^c}[\Gamma_T]$

By hypercontractivity,

$$\mathbb{E}_{T^c}[\Gamma_T] = \mathbb{E}\left[\max_{i \in T} \|\Delta_{i, T^c}\|_2^2\right]$$

$$\leq n \cdot \mathbb{E}\left[\max_{i \in T, j \in T^c} \Delta_{ij}^2\right]$$

$$\leq n \cdot \mathbb{E}\left[\sum_{ij \in [n]} \Delta_{ij}^{2p}\right]^{\frac{1}{p}}$$

$$\leq n^{1 + \frac{2}{p}} \cdot \mathbb{E}[\Delta_{ij}^{2p}]^{\frac{1}{p}}$$

$$\leq n^{1 + \frac{2}{p}} C_p^2 \mathbb{E}[\Delta_{ij}^2]$$

$$\mathbb{E}_{T^c}[\Gamma_T] \leq n^{1 + \frac{2}{p}} C_p^2 \|K\|_{op} \mathcal{H}(K)$$

Putting everything together:

(40)

$$\mathbb{E}[\|\Delta\|_{\text{op}}] \leq K_p \left\{ \|K\|_{\text{op}} \text{Tr}(K) n^{1+\frac{2}{p}} \log n \right\}^{\frac{1}{2}} \\ + K_p \left\{ n \cdot \|K\|_{\text{op}} \mathbb{E}[\|\Delta\|_{\text{op}}] \right\}^{\frac{1}{2}}$$

$$\left[\begin{array}{l} \text{Use that } x^2 - \varepsilon_1 x - \varepsilon_2 \leq 0 \Rightarrow x \leq \frac{\varepsilon_1 + (\varepsilon_1^2 + 4\varepsilon_2)^{\frac{1}{2}}}{2} \\ \leq (\varepsilon_1^2 + 4\varepsilon_2)^{\frac{1}{2}} \end{array} \right]$$

$$\mathbb{E}[\|\Delta\|_{\text{op}}] \leq K_p \left\{ n \|K\|_{\text{op}} + \left[\|K\|_{\text{op}} \text{Tr}(K) n^{1+\frac{2}{p}} \log n \right]^{\frac{1}{2}} \right\}$$

