

MIT 9.520/6.7910
Statistical Learning Theory and Applications

Class 9: Learning bounds for linear least squares

Lorenzo Rosasco

ML algorithms design so far

- ▶ ERM+bias/regularization+optimization
- ▶ Implicit regularization
- ▶ Linear models
- ▶ Nonlinear models: neural nets+kernels

Theory so far

$$\min_{f \in \mathcal{T}} L(f) \quad \mapsto \quad \min_{f \in \mathcal{H}} \widehat{L}(f)$$

► Representation: $\mathcal{H}, \mathcal{B}$ etc.

► Optimization:

$$\min_{f \in \mathcal{H}} \widehat{L}(f).$$

► **Generalization:** $L(\widehat{f})$.

Outline

Recap statistical learning theory

Linear least squares

Error decomposition

Approximation Error

Estimation error

- Estimation Error I: probabilistic estimates

- Estimation Error II: deterministic estimates

Final bound

Ghost track: improved approximation results

The supervised learning problem

- ▶ $\mathcal{X} \times \mathbb{R}$ probability space, with distribution P .
- ▶ $\ell : \mathbb{R} \times \mathbb{R} \rightarrow [0, \infty)$ *loss function*.

Problem: Solve

$$L(f_P) = \min_{f \in \mathcal{T}} L(f), \quad L(f) = \mathbb{E}_{(x,y) \sim P}[\ell(y, f(x))],$$

given

$$S_n = (x_1, y_1), \dots, (x_n, y_n) \sim P^n.$$

ERM

$$S_n \mapsto \hat{f} = \hat{f}_S \in \mathcal{H}$$

ERM

$$\hat{f} = \arg \min_{f \in \mathcal{H}} \hat{L}(f), \quad \hat{L}(f) = \frac{1}{n} \sum_{i=1}^n \ell(y_i, f(x_i)).$$

Note: an estimator is a random variable.

How good is a solution?

$$S_n \rightarrow \hat{f}$$

- Practice: **test error**¹

$$\hat{L}_{\text{Test}}(\hat{f}) = \frac{2}{n} \sum_{i=n/2+1}^n \ell(y_i, \hat{f}(x_i)).$$

- Theory: **excess risk**

$$L(\hat{f}) - L(f_P).$$

Note: these are also random variables.

¹Use $n/2$ points for $\hat{f}$.

Outline

Recap statistical learning theory

Linear least squares

Error decomposition

Approximation Error

Estimation error

- Estimation Error I: probabilistic estimates

- Estimation Error II: deterministic estimates

Final bound

Ghost track: improved approximation results

Linear least squares

- ▶ Linear models

$$f_w(x) = w^\top x,$$

(analysis applies to $f_w(x) = \langle w, \Phi(x) \rangle$).

- ▶ Linear least squares

$$(y - w^\top x)^2.$$

Lots of nice formulas!

- ▶ Bounded data

$$|y| \leq M, \quad \|x\| \leq \kappa \quad \text{a.s.}$$

Notation for the risks

$$f_w(x) = w^\top x$$

Expected risk

$$L(w) = \mathbb{E}[(y - w^\top x)^2].$$

Empirical risk

$$\hat{L}(w) = \frac{1}{n} \sum_{i=1}^n (y_i - w^\top x_i)^2.$$

Rewriting ridge regression

$$\hat{w}_\lambda = \arg \min_{w \in \mathbb{R}^d} \hat{L}(w) + \lambda \|w\|^2$$

Then

$$\hat{w}_\lambda = (\hat{\Sigma} + \lambda I)^{-1} \hat{h}$$

where

$$\hat{\Sigma} = \frac{1}{n} \sum_{i=1}^n x_i x_i^\top, \quad \hat{h} = \frac{1}{n} \sum_{i=1}^n x_i y_i.$$

Best target in the model

$$L(w^\dagger) = \min_{w \in \mathbb{R}^d} L(w)$$

- ▶ Always true with finite features $f(x) = \langle w, \Phi(x) \rangle$, with $\Phi(x) \in \mathbb{R}^p$.
- ▶ It is an assumption for infinite-dimensional models².

²Think of learning a function in a Sobolev space with the Gaussian kernel.

Normal equation

$$L(w^\dagger) = \min_{w \in \mathbb{R}^d} L(w)$$

$$\Sigma w^\dagger = h.$$

where

$$\Sigma = \mathbb{E}[xx^\top], \quad h = \mathbb{E}[xy].$$

We define $w^\dagger$ as the minimum-norm solution.

Normal equation: proof

$$\Sigma w^\dagger = h$$

Indeed,

$$\nabla L(w^\dagger) = 0$$

and

$$\nabla L(w) = 2 \mathbb{E}[x(x^\top w - y)] = 2(\Sigma w - h).$$

The minimal-norm solution is the unique one in $\text{Null}(\Sigma)^\perp$.

Covariance matrix

$$\Sigma = \mathbb{E}[xx^\top].$$

This is the covariance (or second-order) matrix. It is:

- ▶ Symmetric.
- ▶ Positive semidefinite.

Excess risk and norm

$$L(w) - L(w^\dagger) = \|\Sigma^{1/2}(w - w^\dagger)\|^2$$

Proof: adding and subtracting $x^\top w^\dagger$:

$$\begin{aligned} L(w) &= \mathbb{E}[(y - x^\top w^\dagger + x^\top (w^\dagger - w))^2] \\ &= \mathbb{E}[(y - x^\top w^\dagger)^2] + \mathbb{E}[(x^\top (w - w^\dagger))^2] + 2 \mathbb{E}[(y - x^\top w^\dagger) x^\top (w - w^\dagger)]. \end{aligned}$$

First term is $L(w^\dagger)$ and the last term is zero since

$$\mathbb{E}[(y - x^\top w^\dagger) x] = h - \Sigma w^\dagger = 0.$$

Finally,

$$L(w) - L(w^\dagger) = (w - w^\dagger)^\top \Sigma (w - w^\dagger) = \|\Sigma^{1/2}(w - w^\dagger)\|^2.$$

Excess risk of least squares

$$L(\hat{w}_\lambda) - L(w^\dagger) = \|\Sigma^{1/2}(\hat{w}_\lambda - w^\dagger)\|^2.$$

where

$$\hat{w}_\lambda = (\hat{\Sigma} + \lambda I)^{-1} \hat{h}, \quad \Sigma w^\dagger = h.$$

Matrix norm and spectral theorem

Let A be a symmetric positive semidefinite matrix with eigenvalues $\lambda_i(A)$.

Then

$$\|A\| = \max_i \lambda_i(A),$$

and for any continuous $f : [0, \infty) \rightarrow \mathbb{R}$

$$\|f(A)\| = \max_i |f(\lambda_i(A))|.$$

Moreover

$$\|A\| \leq \|A\|_F \leq \text{Tr}(A),$$

where the Frobenius norm is $\|A\|_F^2 = \langle A, A \rangle_F = \text{Tr}(A^\top A)$.

Outline

Recap statistical learning theory

Linear least squares

Error decomposition

Approximation Error

Estimation error

- Estimation Error I: probabilistic estimates

- Estimation Error II: deterministic estimates

Final bound

Ghost track: improved approximation results

Decomposition

$$\|\Sigma^{1/2}(\hat{w}_\lambda - w^\dagger)\|^2$$

$$\hat{w}_\lambda - w^\dagger = \hat{w}_\lambda - w_\lambda + w_\lambda - w^\dagger$$

where

$$w_\lambda = \arg \min_{w \in \mathbb{R}^d} L(w) + \lambda \|w\|^2$$

and

$$\hat{w}_\lambda = \arg \min_{w \in \mathbb{R}^d} \hat{L}(w) + \lambda \|w\|^2, \quad w^\dagger = \arg \min_{w \in \mathbb{R}^d} L(w).$$

Decomposition (cont.)

$$\|\Sigma^{1/2}(\hat{w}_\lambda - w^\dagger)\|^2$$

$$\hat{w}_\lambda - w^\dagger = \hat{w}_\lambda - w_\lambda + w_\lambda - w^\dagger$$

where

$$w_\lambda = (\Sigma + \lambda I)^{-1}h$$

and

$$\hat{w}_\lambda = (\hat{\Sigma} + \lambda I)^{-1}\hat{h} \qquad \Sigma w^\dagger = h.$$

Approximation and estimation errors³

$$\|\Sigma^{1/2}(\hat{w}_\lambda - w^\dagger)\|^2, \quad \hat{w}_\lambda - w^\dagger = \hat{w}_\lambda - w_\lambda + w_\lambda - w^\dagger$$

$$\left\| \Sigma^{1/2}(\hat{w}_\lambda - w^\dagger) \right\|^2 \leq 2 \left\| \Sigma^{1/2}(\hat{w}_\lambda - w_\lambda) \right\|^2 + 2 \left\| \Sigma^{1/2}(w_\lambda - w^\dagger) \right\|^2$$

- Estimation/Variance (stochastic)

$$\left\| \Sigma^{1/2}(\hat{w}_\lambda - w_\lambda) \right\|^2.$$

- Approximation/Bias (deterministic)

$$\left\| \Sigma^{1/2}(w_\lambda - w^\dagger) \right\|^2.$$

³Aka Bias-Variance decomposition

Outline

Recap statistical learning theory

Linear least squares

Error decomposition

Approximation Error

Estimation error

Estimation Error I: probabilistic estimates

Estimation Error II: deterministic estimates

Final bound

Ghost track: improved approximation results

Approximation and estimation errors⁴

$$\left\| \Sigma^{1/2}(\hat{w}_\lambda - w^\dagger) \right\|^2 \leq 2 \left\| \Sigma^{1/2}(\hat{w}_\lambda - w_\lambda) \right\|^2 + 2 \left\| \Sigma^{1/2}(w_\lambda - w^\dagger) \right\|^2$$

- Estimation/Variance (stochastic)

$$\left\| \Sigma^{1/2}(\hat{w}_\lambda - w_\lambda) \right\|^2$$

- **Approximation/Bias** (deterministic)

$$\left\| \Sigma^{1/2}(w_\lambda - w^\dagger) \right\|^2$$

⁴Aka Bias-Variance decomposition

Approximation Error

For $\lambda \leq \kappa$

$$\|\Sigma^{1/2}(w_\lambda - w^\dagger)\|^2 \leq \lambda \|w^\dagger\|^2.$$

Proof: $w_\lambda = (\Sigma + \lambda I)^{-1}h$, $\Sigma w^\dagger = h$, so $w_\lambda = (\Sigma + \lambda I)^{-1}\Sigma w^\dagger$. Then⁵

$$w_\lambda - w^\dagger = ((\Sigma + \lambda I)^{-1}\Sigma - I)w^\dagger = (\Sigma + \lambda I)^{-1}(\Sigma - (\Sigma + \lambda I))w^\dagger = -\lambda(\Sigma + \lambda I)^{-1}w^\dagger.$$

Hence

$$\|\Sigma^{1/2}(w_\lambda - w^\dagger)\| \leq \lambda \|\Sigma^{1/2}(\Sigma + \lambda I)^{-1}\| \|w^\dagger\|$$

Since

$$\|\Sigma^{1/2}(\Sigma + \lambda I)^{-1}\| = \max_i \frac{\sqrt{\lambda_i}}{\lambda_i + \lambda} = \frac{1}{2\sqrt{\lambda}} \quad (\lambda \leq \kappa),$$

then

$$\|\Sigma^{1/2}(w_\lambda - w^\dagger)\| \leq \sqrt{\lambda} \|w^\dagger\|.$$

⁵ $I = (\Sigma + \lambda I)^{-1}(\Sigma + \lambda I)$.

Remarks

$$\|\Sigma^{1/2}(w_\lambda - w^\dagger)\|^2 \leq \lambda \|w^\dagger\|^2.$$

- ▶ The result is dimension free and sharp.
- ▶ It can be improved if we allow a dependence on $\lambda_{\min}(\Sigma)$ or add assumptions.

Improved approximation error

$$\|\Sigma^{1/2}(w_\lambda - w^\dagger)\|^2 \leq \lambda \|w^\dagger\|^2.$$

- It can be improved if we allow a dependence on⁶ $\lambda_{\min}(\Sigma)$. If $\lambda \leq \lambda_{\min}(\Sigma)$, then

$$\|\Sigma^{1/2}(\Sigma + \lambda I)^{-1}\| \leq \frac{\sqrt{\lambda_{\min}(\Sigma)}}{\lambda_{\min}(\Sigma) + \lambda} \quad \Rightarrow \quad \|\Sigma^{1/2}(w_\lambda - w^\dagger)\| \leq \frac{\lambda}{\sqrt{\lambda_{\min}(\Sigma)}} \|w^\dagger\|.$$

- It can also be improved if $w^\dagger$ is *aligned* with Σ ,

$$w^\dagger = \Sigma^\alpha v^\dagger \quad \Rightarrow \quad \|\Sigma^{1/2}(w_\lambda - w^\dagger)\| \leq \lambda^{\alpha+1/2} \|v^\dagger\|.$$

This is called the source condition.

⁶ $\lambda_{\min}(\Sigma)$ is always zero for infinite dimensional features!

Outline

Recap statistical learning theory

Linear least squares

Error decomposition

Approximation Error

Estimation error

Estimation Error I: probabilistic estimates

Estimation Error II: deterministic estimates

Final bound

Ghost track: improved approximation results

Approximation and estimation errors

$$\left\| \Sigma^{1/2}(\hat{w}_\lambda - w^\dagger) \right\|^2 \leq 2 \left\| \Sigma^{1/2}(\hat{w}_\lambda - w_\lambda) \right\|^2 + 2 \left\| \Sigma^{1/2}(w_\lambda - w^\dagger) \right\|^2$$

► Estimation/Variance (stochastic)

$$\left\| \Sigma^{1/2}(\hat{w}_\lambda - w_\lambda) \right\|^2$$

► Approximation/Bias (deterministic) ✓

$$\left\| \Sigma^{1/2}(w_\lambda - w^\dagger) \right\|^2 \leq \lambda \left\| w^\dagger \right\|^2.$$

Estimation error

For $\lambda_n \leq \lambda \leq \kappa$, w.h.p.⁷

$$\left\| \Sigma^{1/2}(\hat{w}_\lambda - w_\lambda) \right\|^2 \leq \frac{C}{\lambda n}$$

Key idea:

$$\hat{w}_\lambda = (\hat{\Sigma} + \lambda I)^{-1} \hat{h}, \quad w_\lambda = (\Sigma + \lambda I)^{-1} h$$

should be close if

$$\hat{h} \approx h \quad \text{and} \quad \hat{\Sigma} \approx \Sigma.$$

⁷With high probability.

Hoeffding inequality: scalar case

Let $Z_1, \dots, Z_n \in \mathbb{R}$ be independent random variables with

$$|Z_i| \leq C \quad \text{a.s.},$$

and $\mu = \mathbb{E}[Z_i]$. Then, for any $\epsilon > 0$,

$$\Pr\left(\left|\frac{1}{n} \sum_{i=1}^n Z_i - \mu\right| \geq \epsilon\right) \leq 2 \exp\left(-\frac{n\epsilon^2}{2C^2}\right).$$

Equivalently, with probability at least $1 - \delta$,

$$\left|\frac{1}{n} \sum_{i=1}^n Z_i - \mu\right| \leq C \sqrt{\frac{2 \log(2/\delta)}{n}}.$$

Hoeffding inequality: vector case

Let $Z_1, \dots, Z_n \in \mathcal{H}$ be independent random variables in a Hilbert space $(\mathcal{H}, \langle \cdot, \cdot \rangle)$ with

$$\|Z_i\| \leq C \quad \text{a.s.,}$$

and $\mu = \mathbb{E}[Z_i]$. Then, for any $\epsilon > 0$,

$$\Pr\left(\left\|\frac{1}{n} \sum_{i=1}^n Z_i - \mu\right\| \geq \epsilon\right) \leq 2 \exp\left(-\frac{n\epsilon^2}{2C^2}\right).$$

Equivalently, with probability at least $1 - \delta$,

$$\left\|\frac{1}{n} \sum_{i=1}^n Z_i - \mu\right\| \leq C \sqrt{\frac{2 \log(2/\delta)}{n}}.$$

Probabilistic estimates

With probability $1 - \delta$

$$\left\| \hat{h} - h \right\| \leq \kappa M \sqrt{\frac{2 \log(4/\delta)}{n}}, \quad \text{and} \quad \left\| \hat{\Sigma} - \Sigma \right\|_{\text{F}} \leq \kappa^2 \sqrt{\frac{2 \log(4/\delta)}{n}}.$$

- ▶ $Z_i = x_i y_i$, $\mu = h$, and $\|Z_i\| \leq \kappa M$.
- ▶ $Z_i = x_i x_i^\top$, $\mu = \Sigma$, and⁸ $\|Z_i\|_{\text{F}} \leq \kappa^2$
- ▶ Union bound to control both events simultaneously ($\delta \mapsto \delta/2$).

⁸ $\|x_i x_i^\top\|_{\text{F}} \leq \text{Tr}(x_i x_i^\top) = \|x_i\|^2 \leq \kappa^2$.

Key deterministic estimate

$$\hat{w}_\lambda - w_\lambda = (\hat{\Sigma} + \lambda I)^{-1} \left[(\hat{h} - h) + (\Sigma - \hat{\Sigma})w_\lambda \right]$$

By definition

$$\hat{w}_\lambda - w_\lambda = (\hat{\Sigma} + \lambda I)^{-1} \hat{h} - (\Sigma + \lambda I)^{-1} h$$

Add and subtract $(\hat{\Sigma} + \lambda I)^{-1} h$

$$\hat{w}_\lambda - w_\lambda = (\hat{\Sigma} + \lambda I)^{-1} (\hat{h} - h) + \left[(\hat{\Sigma} + \lambda I)^{-1} - (\Sigma + \lambda I)^{-1} \right] h$$

Use $A^{-1} - B^{-1} = A^{-1}(B - A)B^{-1}$

$$\hat{w}_\lambda - w_\lambda = (\hat{\Sigma} + \lambda I)^{-1} (\hat{h} - h) + (\hat{\Sigma} + \lambda I)^{-1} (\Sigma - \hat{\Sigma}) \underbrace{(\Sigma + \lambda I)^{-1} h}_{w_\lambda}.$$

Intermediate bound

$$\hat{w}_\lambda - w_\lambda = (\hat{\Sigma} + \lambda I)^{-1} [(\hat{h} - h) + (\Sigma - \hat{\Sigma})w_\lambda]$$

Then

$$\left\| \Sigma^{1/2}(\hat{w}_\lambda - w_\lambda) \right\| \leq \left\| \Sigma^{1/2}(\hat{\Sigma} + \lambda I)^{-1} \right\| \left[\left\| \hat{h} - h \right\| + \left\| \hat{\Sigma} - \Sigma \right\| \|w_\lambda\| \right]$$

► $\|w_\lambda\| \leq \|w^\dagger\|$, since

$$\|w_\lambda\| = \|(\Sigma + \lambda I)^{-1}h\| = \|(\Sigma + \lambda I)^{-1}\Sigma w^\dagger\| \leq \|w^\dagger\|.$$

► The critical term is

$$\left\| \Sigma^{1/2}(\hat{\Sigma} + \lambda I)^{-1} \right\|.$$

Key deterministic estimate II

If $2\kappa^2 \sqrt{\frac{2 \log(4/\delta)}{n}} \leq \lambda \leq \kappa$, then

$$\left\| \Sigma^{1/2} (\hat{\Sigma} + \lambda I)^{-1} \right\| \leq \frac{2}{\sqrt{\lambda}}.$$

Indeed,

$$(\hat{\Sigma} + \lambda I)^{-1} = \left[(\Sigma + \lambda I) + (\hat{\Sigma} - \Sigma) \right]^{-1} = (\Sigma + \lambda I)^{-1} \left[I - (\Sigma - \hat{\Sigma})(\Sigma + \lambda I)^{-1} \right]^{-1}.$$

Using the Neumann series,

$$\left\| \left(I - (\Sigma - \hat{\Sigma})(\Sigma + \lambda I)^{-1} \right)^{-1} \right\| \leq \left(1 - \left\| (\Sigma - \hat{\Sigma})(\Sigma + \lambda I)^{-1} \right\| \right)^{-1}.$$

Finally

$$\left\| \Sigma^{1/2} (\hat{\Sigma} + \lambda I)^{-1} \right\| \leq \underbrace{\left\| \Sigma^{1/2} (\Sigma + \lambda I)^{-1} \right\|}_{\leq \frac{1}{\sqrt{\lambda}}} \underbrace{\left(1 - \left\| (\Sigma - \hat{\Sigma})(\Sigma + \lambda I)^{-1} \right\| \right)^{-1}}_{\leq \frac{\|\hat{\Sigma} - \Sigma\|}{\lambda} \leq \frac{1}{2}}.$$

Neumann series argument

Let A be invertible and set $B = A - C$. Assume $\|CA^{-1}\| < 1$. Then

$$I - CA^{-1} \text{ is invertible, } (I - CA^{-1})^{-1} = \sum_{k=0}^{\infty} (CA^{-1})^k.$$

Hence B is invertible and

$$B^{-1} = A^{-1}(I - CA^{-1})^{-1}, \quad \|B^{-1}\| \leq \frac{\|A^{-1}\|}{1 - \|CA^{-1}\|}.$$

In our case:

$$A = \Sigma + \lambda I, \quad B = \hat{\Sigma} + \lambda I, \quad C = \Sigma - \hat{\Sigma}.$$

Final estimation error bound

With probability $1 - \delta$, if $2\kappa^2 \sqrt{\frac{2 \log(4/\delta)}{n}} \leq \lambda \leq \kappa$, then

$$\begin{aligned} \|\Sigma^{1/2}(\hat{w}_\lambda - w_\lambda)\| &\leq \underbrace{\left\| \Sigma^{1/2}(\hat{\Sigma} + \lambda I)^{-1} \right\|}_{\leq \frac{2}{\sqrt{\lambda}}} \left[\underbrace{\|\hat{h} - h\|}_{\leq \kappa M \sqrt{\frac{2 \log(4/\delta)}{n}}} + \underbrace{\|\hat{\Sigma} - \Sigma\|}_{\leq \kappa^2 \sqrt{\frac{2 \log(4/\delta)}{n}}} \|w^\dagger\| \right] \\ &\leq C \frac{\sqrt{\log(4/\delta)}}{\sqrt{\lambda n}}. \end{aligned}$$

where $C = (2\sqrt{2}(M\kappa + \kappa^2 \|w^\dagger\|))$.

Approximation and estimation errors⁹

$$\left\| \Sigma^{1/2}(\hat{w}_\lambda - w^\dagger) \right\|^2 \leq 2 \left\| \Sigma^{1/2}(\hat{w}_\lambda - w_\lambda) \right\|^2 + 2 \left\| \Sigma^{1/2}(w_\lambda - w^\dagger) \right\|^2$$

- Estimation/Variance (stochastic) ✓

$$\left\| \Sigma^{1/2}(\hat{w}_\lambda - w_\lambda) \right\|^2 \leq C \frac{\log(4/\delta)}{\lambda n}.$$

- Approximation/Bias (deterministic) ✓

$$\left\| \Sigma^{1/2}(w_\lambda - w^\dagger) \right\|^2 \leq \lambda \|w^\dagger\|^2.$$

⁹Aka Bias-Variance decomposition

Outline

Recap statistical learning theory

Linear least squares

Error decomposition

Approximation Error

Estimation error

- Estimation Error I: probabilistic estimates

- Estimation Error II: deterministic estimates

Final bound

Ghost track: improved approximation results

Error bound for fixed λ

If $2\kappa^2 \sqrt{\frac{2 \log(4/\delta)}{n}} \leq \lambda \leq \kappa$, then

$$\left\| \Sigma^{1/2}(\hat{w}_\lambda - w^\dagger) \right\|^2 \leq C^2 \frac{\log(4/\delta)}{\lambda n} + 2\lambda \|w^\dagger\|$$

where $C = 2\sqrt{2} (M\kappa + \kappa^2 \|w^\dagger\|)$.

- ▶ We can choose λ balancing these terms.
- ▶ We can obtain a bound $O(\frac{1}{\sqrt{n}})$.
- ▶ This choice of λ depends on quantities not known in practice, hold-out can be used instead.

Remarks

- ▶ The result is optimal in the considered setting.
- ▶ Better (fast) rates can be obtained with further assumptions (noise, smoothness eigendecay etc.).
- ▶ The bound is valid only for

$$2\kappa^2 \sqrt{\frac{2 \log(2/\delta)}{n}} \leq \lambda$$

recent studies extend this analysis beyond this range (interpolation/benign overfitting).

Summary

- ▶ Generalization error for linear least squares.
- ▶ Dimension-free analysis applying to infinite-dimensional features/kernels.
- ▶ Only one of many possible regimes (faster/slower rates possible!).
- ▶ Many extensions (implicit regularization, benign overfitting, etc.).

Outline

Recap statistical learning theory

Linear least squares

Error decomposition

Approximation Error

Estimation error

- Estimation Error I: probabilistic estimates

- Estimation Error II: deterministic estimates

Final bound

Ghost track: improved approximation results

Improved approximation error

$$\|\Sigma^{1/2}(w_\lambda - w^\dagger)\|^2 \leq \lambda \|w^\dagger\|^2.$$

- ▶ The result is dimension free and sharp.
- ▶ It can be improved if we allow a dependence on¹⁰ $\lambda_{\min}(\Sigma)$. If $\lambda \leq \lambda_{\min}(\Sigma)$, then

$$\|\Sigma^{1/2}(\Sigma + \lambda I)^{-1}\| \leq \frac{\sqrt{\lambda_{\min}(\Sigma)}}{\lambda_{\min}(\Sigma) + \lambda} \Rightarrow \|\Sigma^{1/2}(w_\lambda - w^\dagger)\| \leq \frac{\lambda}{\sqrt{\lambda_{\min}(\Sigma)}} \|w^\dagger\|.$$

- ▶ It can also be improved if $w^\dagger$ is *aligned* with Σ ,

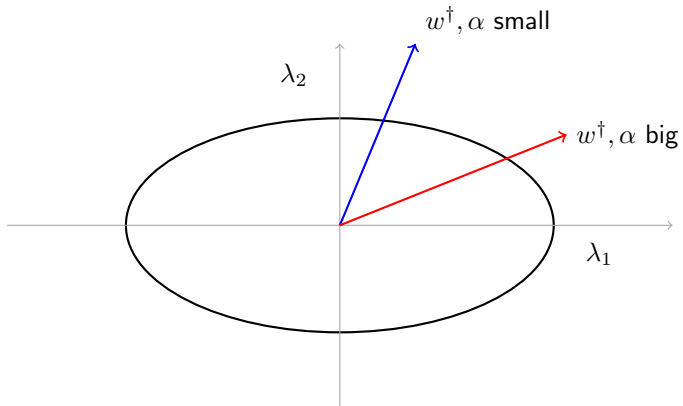
$$w^\dagger = \Sigma^\alpha v^\dagger \Rightarrow \|\Sigma^{1/2}(w_\lambda - w^\dagger)\| \leq \lambda^{\alpha+1/2} \|v^\dagger\|.$$

This is called the source condition.

¹⁰ $\lambda_{\min}(\Sigma)$ is always zero for infinite dimensional features!

Alignment geometry

$$w^\dagger = \sum^\alpha v^\dagger$$



Alignment as smoothness

$$w^\dagger = \Sigma^\alpha v^\dagger$$

- The above condition is equivalent to

$$v^\dagger = \Sigma^{-\alpha} w^\dagger = \sum_i \frac{(v_i^\top w^\dagger)}{\lambda_i^\alpha} v_i \quad \Leftrightarrow \quad \sum_i \frac{(v_i^\top w^\dagger)^2}{\lambda_i^{2\alpha}} < \infty$$

- For feature maps of translation/rotation invariant kernels, the source condition corresponds to decay in the Fourier domain.