

Lecture Notes September 17

OLS, Ridge & Lasso

$$\text{OLS : } \begin{aligned} & \|y - X\beta\|_2^2 \\ & \rightarrow = \frac{1}{n} \sum_{i=1}^{n-1} (y_i - x_i^T \beta)(y_i - x_i^T \beta) \end{aligned}$$

$$\text{Ridge : } \|y - X\beta\|_2^2 + \lambda \| \beta \|_2^2$$
$$\| \beta \|_2^2 = \sum_{i=0}^{p-1} \beta_i^2 \leq t$$

Lasso Regression:

$$\|y - X\beta\|_2^2 + \lambda \| \beta \|_1$$
$$\| \beta \|_1 = \sum_{i=0}^{p-1} |\beta_i| \leq t$$

$$\lambda \geq 0$$

$$\hat{\beta}^{\text{OLS}} = (X^T X)^{-1} X^T y$$

$$\hat{\beta}^{\text{Ridge}} = (X^T X + \lambda I)^{-1} X^T y$$

$\hat{\beta}^{\text{Lasso}}$: not analytical solution

$$\frac{d|x|}{dx} = \text{sign}(x) = \begin{cases} 1 & x > 0 \\ 0 & x = 0 \end{cases}$$

$$(-1 \text{ if } x < 0$$

$$\frac{\partial C^{\text{Lasso}}(\beta)}{\partial \beta} = -2X^T(y - X\beta) + \lambda \text{sign}(\beta)$$

$$\frac{\partial C^{\text{Ridge}}}{\partial \beta} = -2X^T(y - X\beta) + \lambda 2 \cdot \beta$$

$$\frac{\partial C^{\text{OLS}}}{\partial \beta} = -2X^T(y - X\beta)$$

$$\|v\|_2 = \sqrt{\sum_{i=1}^n v_i^2}$$

$$\|v\|_2^2 = \sum v_i^2$$

$$\|v\|_1 = \sum_i |v_i|$$

$$\frac{\partial^2 C^{\text{OLS}}}{\partial \beta \partial \beta^T} = 2X^T X = 2 \cdot \underbrace{(H)}_{\text{Hessian}}$$

$$\text{var}(\hat{\beta}^{\text{OLS}}) = \sigma^2 (X^T X)^{-1}$$

$$\text{var}(\hat{\beta}_i^{\text{OLS}}) = \sigma^2 (X^T X)^{-1}_{ii}$$

$$\varepsilon \sim N(0, \sigma^2)$$

$$\text{covariance matrix} = C[x]$$

$$= \frac{1}{n} X^T X$$

$$= \begin{bmatrix} \text{var}[x_0] & \text{cov}[x_0 x_1] & \dots & \text{cov}[x_0 x_{p-1}] \\ \vdots & & & \\ \text{cov}[x_{p-1} x_0] & \dots & \dots & \text{var}[x_{p-1}] \end{bmatrix},$$

$$\text{var}[x_0] = \frac{1}{n} \sum_{i=0}^{n-1} (x_{i0} - \bar{x}_0)^2$$

$\hat{H} \propto C[x]$ a positive definite matrix

$\Rightarrow \hat{H} > 0 \Rightarrow \hat{\beta}$ is a true minimum of $C(\beta)$

$$\frac{\partial^2 C^{\text{Ridge}}}{\partial \beta \partial \beta^T} = 2 \boxed{H} + 2\lambda I$$

$$\frac{\partial^2 C^{\text{Lasso}}}{\partial \beta \partial \beta^T} = 2H$$

$$\frac{\partial C^{\text{lasso}}}{\partial \beta} = -2X^T y + 2H\beta + \lambda \text{sign}(\beta)$$

$C[X] = \text{covariance matrix}$

$$C(\beta) = C(y|X\beta) = \text{cost function}$$

$$= L(y|X\beta) \\ \text{loss function}$$

$$= R(y|X\beta) \\ \text{Risk function,}$$

$$\begin{cases} \text{var}[x] & E[x] \\ \text{cov}[x, y] \end{cases} \quad \text{statistical properties.}$$

$$(X^T X) V = V \Sigma^2$$

$$X = U \Sigma V^T$$

$$X \in \mathbb{R}^{n \times p}$$

$$U \in \mathbb{R}^{n \times n}$$

$$\Sigma \in \mathbb{R}^{n \times p}$$

$$V \in \mathbb{R}^{p \times p}$$

$$UU^T = U^T U = \mathbb{1}$$

$$V^T V = V V^T = \mathbb{1}$$

$$(X X^T) U = U \Sigma \Sigma^T = U \Sigma^2$$

$$\hat{y}^{OLS} = X \hat{\beta}^{OLS} = \sum_{i=0}^{p-1} u_i u_i^T y$$

$$\hat{y}^{Ridge} = X \hat{\beta}^{Ridge} = \sum_{i=0}^{p-1} u_i \frac{\sigma_i^2}{\sigma_i^2 + \lambda} u_i^T y$$

$$\sigma_i \geq \sigma_{i+1}$$

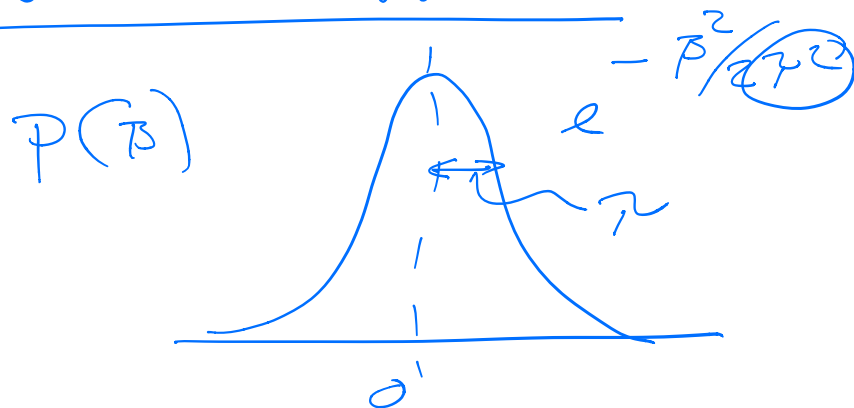
Ridge shrinks the singular values, reducing more (for a given λ) those σ_i which are smaller. These correspond to eigenvalues with small variance.

$$\begin{aligned}
 \beta^{OLS} &= \frac{(X^T X)^{-1} X^T y}{1} \\
 &= (V \Sigma^T U^T U \Sigma V^T)^{-1} V \Sigma^T U^T y \\
 &= (\underline{\Sigma^T \Sigma})^{-1} \underline{V \Sigma^T U^T} y
 \end{aligned}$$

$$\begin{aligned}
 \text{var}[\hat{\beta}^{\text{Ridge}}] &= \sigma^2 [X^T X + \lambda I]^{-1} \\
 &\quad \times X^T X [\boxed{X^T X} + \lambda I]
 \end{aligned}$$

Ridge shrinks the variance of $\hat{\beta}_i^{\text{Ridge}}$ more for those β_i values with small singular values.

Bayesian approach



$$X^T X = \underline{1} = \begin{bmatrix} 1 & 0 \\ 0 & \ddots \\ & & 1 \end{bmatrix}$$

$$\frac{\hat{y}^{\text{Ridge}}}{\frac{1}{1+\lambda}} = \hat{y}^{\text{OLS}} =$$

$$\frac{X^T X = \underline{1}}{X^T X \neq \underline{1}}$$

$$\frac{1}{1+\lambda} y$$

$$\begin{aligned} \text{var}[\hat{\beta}^{\text{Ridge}}] &= \sigma^2 \frac{1}{\mathbb{I}(1+\lambda)} \mathbb{I} \frac{1}{\mathbb{I}(1+\lambda)} \\ &= \sigma^2 \frac{1}{(1+\lambda)^2} \end{aligned}$$

————— * —————

Degrees of freedom

p - predictors/features . . .

$$\hat{y}^{\text{OLS}} = \sum_{i=0}^{p-1} u_i u_i^T y$$

$$= X(X^T X)^{-1} X^T y$$

$$\text{Tr}[X(X^T X)^{-1} X^T] = p$$

= degrees of freedom,

$$\hat{y}^{\text{Ridge}} = X(X^T X^{-1} + \lambda \mathbb{I})^{-1} X^T y$$

$$\text{tr} \left[X (X^T X + \lambda I)^{-1} X^T \right] = \sum_{j=0}^{p-1} \frac{\sigma_j^2}{\sigma_j^2 + \lambda} = df(\lambda)$$

$$\lambda = 0 \Rightarrow df(\lambda) = p$$

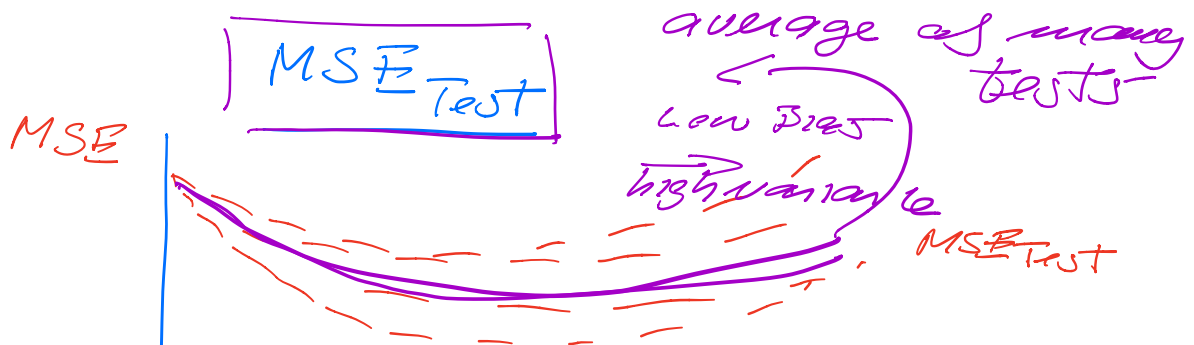
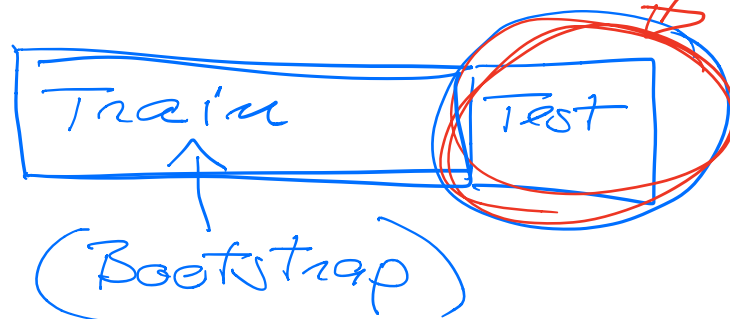
$$\text{tr} [A] = \sum_{i=1}^n a_{ii}$$

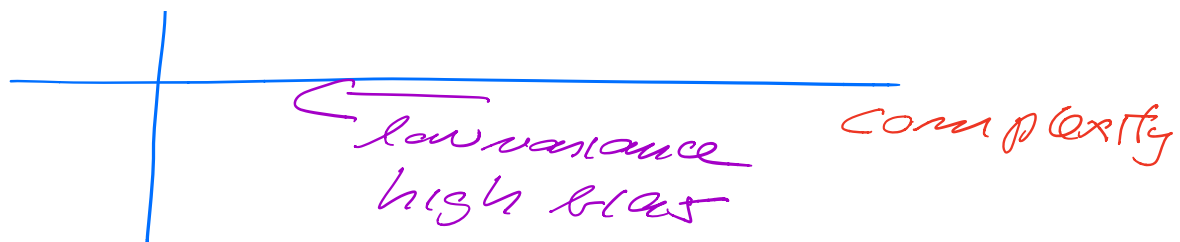
$$\lambda \rightarrow \infty \Rightarrow df(\lambda) = 0$$

$$\hat{y} = X\beta$$

Resampling Methods

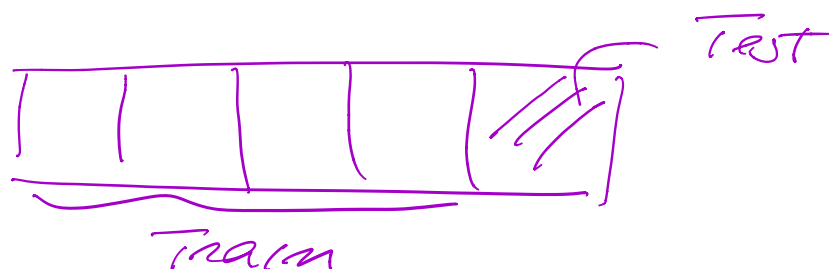
1b : Bootstrap + bias-variance (OLS + Ridge) + 10x





1C OLS + cross-validation

K-Folds $K=5$

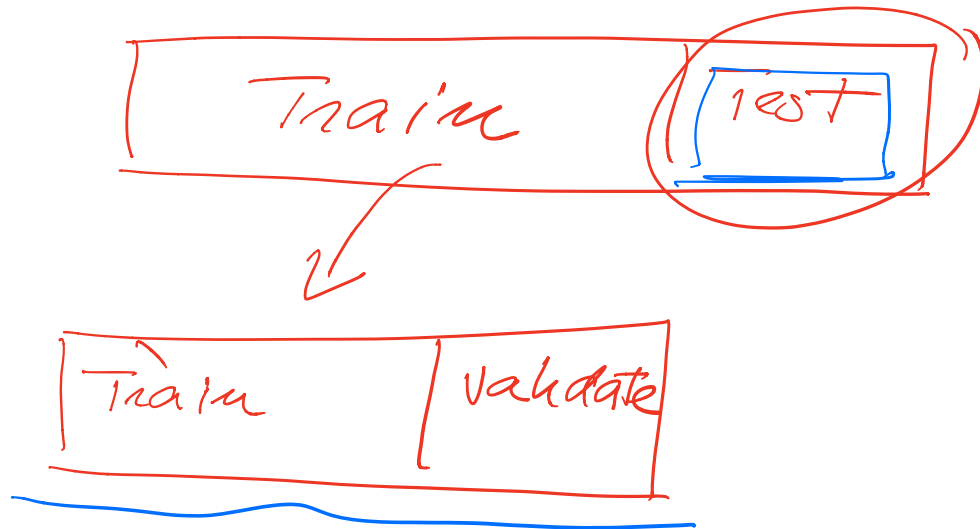


3 more times

in total 5 MSE_{Test}
can be calculated.

— No Need to do a Train-Test split.

Train-test split



Cross-validate here,

Ridge : polynomial
complexity |
 λ -parameter,

validate gives optimal

$MSE(\lambda, \text{polynomial complexity})$

Final model

used on the independent
test data.