

Lecture FYS- STK3155/4155, August 26, 2024

FYS-STK 3155/4155 A06 26

$$C(\beta) = \text{MSE} = \frac{1}{n} \sum_{i=0}^{n-1} \left(y_i - \sum_{j=0}^{p-1} \beta_j x_i^j \right)^2$$

$$\hat{y}_i = \sum_{j=0}^{p-1} \beta_j x_i^j \quad (\text{Model})$$

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

↑
optimal

X orthogonal

$$X^T X = X X^T = I$$

$$= \begin{bmatrix} 1 & & 0 \\ & \ddots & \\ 0 & & 1 \end{bmatrix}$$

$$\hat{y} = X \hat{\beta} = X \underbrace{(X^T X)^{-1}}_I X^T y$$

$$= y$$

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

$$A^2 = A \quad (\text{Projection matrix})$$

$$\tilde{y} = Ay$$

$\tilde{y}$ can be viewed as a projection of y in terms of the column vectors of X

$$X = \begin{bmatrix} 1 & x_0 & x_0^2 & \dots & x_0^{p-1} \\ 1 & x_1 & & & \\ \vdots & \vdots & & & \\ 1 & x_{n-1} & \dots & x_{n-1}^{p-1} \end{bmatrix}$$

$$\frac{\partial C}{\partial \beta} = 0 = -\frac{2}{n} X^T (y - X\beta)$$

Ex 1

$$\frac{\partial^2 C}{\partial \beta \partial \beta^T} = \frac{2}{n} \underbrace{X^T X}_{\text{Hessian}}$$

(MSE)

SVD of a matrix

Decompose Matrix X

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

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

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

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

$$U U^T = U^T U = I$$

$$V V^T = V^T V = I$$

$$U = [u_0 \ u_1 \ \dots \ u_{n-1}]$$

$$u_i^\top u_j = \delta_{ij}$$

$$V = [v_0 \ v_1 \ \dots \ v_{p-1}]$$

$$v_i^\top v_j = \delta_{ij}$$

$$\Sigma \in \mathbb{R}^{n \times p} = \begin{bmatrix} \sigma_0 & & & 0 \\ & \sigma_1 & & \\ & & \ddots & \\ 0 & & & \sigma_{p-1} \\ & & & & 0 \end{bmatrix}$$

$\sigma_0 > \sigma_1 > \sigma_2 > \dots > \sigma_{p-1} > 0$

Example 3×2

$$\Sigma = \begin{bmatrix} 2 & 0 \\ 0 & 1 \\ 0 & 0 \end{bmatrix}$$

$$\Sigma^T \Sigma = \begin{bmatrix} 4 & 0 \\ 0 & 1 \end{bmatrix} = \begin{bmatrix} \sigma_0^2 & 0 \\ 0 & \sigma_1^2 \end{bmatrix}$$

$$\Sigma \Sigma^T = \begin{bmatrix} 4 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0 \end{bmatrix}$$

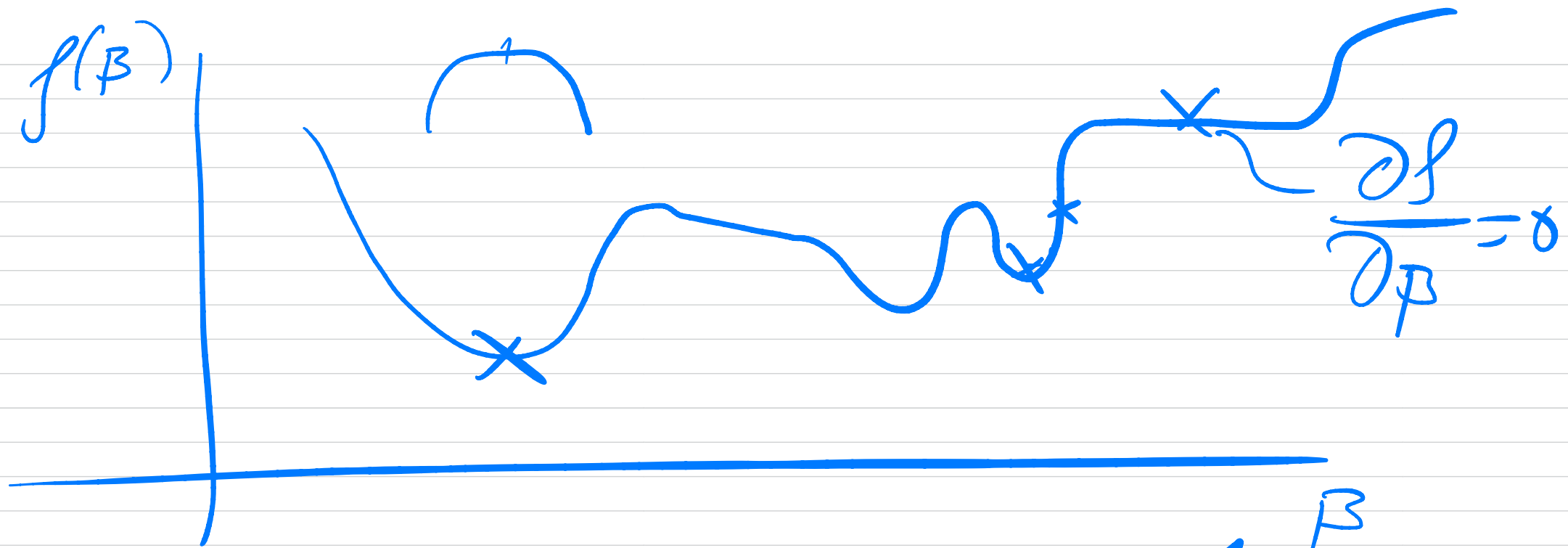
2nd order only

$$\underbrace{X^T X}_{p \times p} = \underbrace{V}_{p \times n} \underbrace{\Sigma^T U^T U \Sigma}_{n \times n} \underbrace{V^T}_{n \times p} \quad p \times p \quad p \times n \quad 1_{n \times n} \quad n \times p \quad p \times p$$

$$= V \Sigma^T \Sigma V^T$$

$$X^T X V = V \Sigma^T \Sigma \underbrace{V^T V}_{I_{p \times p}}$$

$$X^T X v_i = v_i \lambda_i^2 \Rightarrow \text{convex optimization}$$



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

with SVD

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

$$\underset{\substack{\uparrow \\ n}}{\tilde{y}_{OLS}} = \underbrace{U \Sigma V^T}_{n \times p} \left(\underbrace{V \Sigma^T \Sigma V^T}_{p \times p} \right)^{-1} \underbrace{V \Sigma^T U^T}_{p \times n} y$$

$$\underbrace{V \Sigma^T \Sigma V^T}$$

A, B are
square invertible
matrices

$$\frac{1}{AB} = A^{-1} B^{-1}$$

$$\tilde{y} = \left(\sum_{j=0}^{p-1} u_j a_j^T \right) y$$

$$\text{if } \forall \Delta_i > 0$$

$$\tilde{y} = \underbrace{\sum_{j=0}^{n-1} u_j a_j^T}_{\uparrow} y$$

$$\tilde{y} = y$$

why interesting?
our Next Method:
Ridge Regression.

Trick in case you cannot
invert $X^T X$

$$X^T X \rightarrow X^T X + \lambda I_{p \times p}$$

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

$\lambda > 0$

Another cost function

$$C_{\text{Ridge}} = \frac{1}{n} \sum_{i=0}^{n-1} (y_i - \sum_{j=0}^{p-1} \beta_j x_i^j)^2$$

$$+ \lambda \sum_{j=0}^{p-1} \beta_j^2$$



Hyperparameter

constraint

$$\sum_{j=0}^{p-1} \beta_j^2 < t < \infty$$

$$C_{\text{LASSO}} = \frac{1}{n} \sum_{i=0}^{n-1} \left(y_i - \sum_{j=0}^{p-1} \beta_j x_i^j \right)^2$$

$$+ \underbrace{\lambda \sum_{j=0}^{p-1} |\beta_j|}_{\lambda \|\beta\|_1}$$

$$\frac{\partial C_{\text{ridge}}}{\partial \beta} = 0 = -\frac{2}{n} X^T (y - X\beta)$$

$$\rightarrow 2\lambda \beta$$

$$-X^T(Y - X\beta) - \underbrace{n\lambda}_{\lambda \rightarrow \lambda} \beta$$

$$\beta_{\text{ridge}} = \left(X^T X + \lambda I \right)^{-1} X^T Y$$

$p \times p$

Lasso

$$\frac{d|\beta|}{d\beta} = \begin{cases} +1 & \beta > 0 \\ -1 & \beta < 0 \end{cases}$$

$$\hat{y}_{\text{Ridge}} = X \left(\underbrace{X^T X + \lambda I}_{u \Sigma v^T} \right)^{-1} X^T y$$

SVD

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

$$\sigma_0^2 > \sigma_1^2 > \dots > \sigma_{p-1}^2 > 0$$

with Δ_j small and λ_j , we can suppress (reduce) the role of a specific component

$u_j \Rightarrow$

reduction of degrees of freedom.

