

# <Statistics 1>

## 1. Supervised Learning.

### 1.1. Basic Model

input:  $x \in \mathbb{R}^d \Rightarrow x = (x_1, \dots, x_d)$

output:  $y \in \mathcal{Y}$  classification  
regression

model:  $y \approx f(x)$ ,  $f \in \mathcal{M}$ .

$$y = f(x) + \epsilon. \quad \text{additive} \quad y = f(x) + \epsilon$$

loss function:  $\ell(y, a) \begin{cases} \ell(y, a) = (y-a)^2 \\ \ell(y, a) = 1 (y \neq a). \end{cases}$

goal:  $f^0 = \arg \min_{f \in \mathcal{F}} \mathbb{E}_{(x,y) \sim P} [\ell(Y, f(x))]$ . ( $\mathcal{F} \subseteq \mathcal{M}$ ).

prediction:  $f^0$  is unknown.  $\rightarrow$  estimate  $f^0$  by  $\hat{f}$ .

$$\begin{aligned} \hat{f} &= \arg \min_{f \in \mathcal{F}} \mathbb{E}_P [\ell(Y, f(x))] & P_n: \text{empirical distribution} \\ &= \arg \min_f \frac{1}{n} \sum_{i=1}^n \ell(y_i, f(x_i)) & P_n: \frac{1}{n} \sum_{i=1}^n \delta_{(x_i, y_i)} \end{aligned}$$

### 1.2. Linear Regression and Least Neighbors.

Linear Reg. Model:  $\mathcal{M} = \mathcal{F} = \left\{ f_0 + \sum_{j=1}^d \beta_j x_j : \beta_j \in \mathbb{R} \right\}$ .

$\rightarrow$  use least squares: empirical loss = RSS.

$$\begin{aligned} \text{RSS}(f_0) &= \sum_{i=1}^n (y_i - f(x_i))^2 \\ &= \sum_{i=1}^n (y_i - f_0 - \sum_{j=1}^d \beta_j x_{ij})^2 \end{aligned}$$

k-Nearest Neighbor.

$$\hat{f}(x) = \frac{1}{k} \sum_{x_i \in N_k(x)} y_i. \quad N_k(x): \text{neighborhood of } x.$$

### 1.3. Statistical Decision Theory.

Risk / Expected Prediction Error (EPE)

$$\begin{aligned} R(f) &= \mathbb{E}_{(x,y) \sim P} [\ell(Y, f(x))] = \mathbb{E}[(Y - f(x))^2] \\ &= \mathbb{E}_x [\mathbb{E}_{Y|X} [(Y - f(x))^2 | X]] \end{aligned}$$

$$\text{where } \mathbb{E}_{Y|X} [(Y - f(x))^2 | X] = \mathbb{E}_{Y|X} [(Y - \mathbb{E}[Y|X])^2 | X] + (\mathbb{E}[Y|X] - f(x))^2$$

$\Rightarrow R(f)$  is minimized when  $f^0(x) = \mathbb{E}[Y|X=x]$ . "regression function" (usually unknown.)

k-NN)  $\hat{f}(x) = \text{Ave}(y_i | x_i \in N_k(x))$ .

$\rightarrow$  two approximations expectation  $\rightarrow$  average.

conditioning at a point  $\rightarrow$  conditioning on neighbor of  $x$ .

$\rightarrow$  under mild regularity conditions on  $P(X, Y)$ ,  $\hat{f}(x) \rightarrow \mathbb{E}[Y|X=x]$   
as  $n, k \rightarrow \infty$  s.t.  $\frac{k}{n} \rightarrow 0$ . model complexity grows slower than sample size. (asymptotic sample efficiency).

— linear reg.) assume  $f^*(x)$  linear.  $\mathbb{E}[Y|X=x] = x^T \beta$ .

→ model-based approach) let  $\beta^* = (\mathbb{E}[X^T X])^{-1} \mathbb{E}[X^T Y]$ .

$$R(\beta) = \mathbb{E}[(Y - X^T \beta)^2]$$

$$= \mathbb{E}[(Y - X^T \beta^*)^2] + \mathbb{E}[(X^T \beta^* - X^T \beta)^2]. \Rightarrow \beta^* \text{ minimizing the risk.}$$

→  $\hat{\beta} = (X^T X)^{-1} X^T Y$ : average over empirical dist.

→  $\beta \mapsto \beta^*$ . ( $\because$  L.L.N.)

$$\sqrt{n}(\hat{\beta} - \beta^*) \xrightarrow{d} N(0, \Sigma).$$

— classification ( $y = \{0, \dots, K-1\}$ .)

$$R(\theta) = \mathbb{E}_{(Y,X) \sim P} [\ell(Y, f(X))] = \mathbb{E}[\mathbb{I}(Y \neq f(X))].$$

$$= \mathbb{E}_X [\mathbb{E}_Y [\ell(y, f(x)) | X]]$$

$$= \mathbb{E}_X \left[ \sum_{j=0}^{K-1} \ell(j, f(x)) P(Y=j|X) \right].$$

→ minimized when  $f^*(x) = \arg \min_{j=0, \dots, K-1} \sum_{j=0}^{K-1} \ell(j, k) P(Y=j|X=x).$

$$= \arg \max_{j=0, \dots, K-1} P(Y=j|X=x).$$

$\hookrightarrow$  usually unknown.

→ optimal classifier = Bayes classifier.

— k-NN)  $P(Y=j|X=x) \approx \frac{1}{K} \sum_{x_i \in N_K(x)} \mathbb{I}(y_i = j).$

2.4. Curse of Dimensionality.

—  $X = (X_1, \dots, X_d) \sim \text{Uniform}([0,1]^d).$

hypercube neighbor  $\rightarrow$  side length  $e_d(r) = r^{1/d}$ . (capture fraction  $r$  of the obs.)

but  $e_0(0.01) = 0.63$ ,  $e_0(0.1) = 0.50$ .  $\rightarrow$  no longer "local"

—  $X = (X_1, \dots, X_d) \sim \text{Uniform}(B_{R_2}(0,1))$

where  $B_{R_2}(0,1) = \{x \in \mathbb{R}^d : \|x\|_2 \leq 1\}.$

$$X_{i1} = (X_{i1}, \dots, X_{id}), R_2 = \sqrt{\sum_{j=1}^d X_{ij}^2} \rightarrow \text{Med}(R_{(1)}) = \left(1 - \left(\frac{1}{2}\right)^{\frac{1}{d}}\right)^{\frac{1}{2}}.$$

$$\text{if } n=5000 \text{ \& } d=10, \text{ Med}(R_{(1)}) \approx 0.72$$

— (radius of inner sphere) = (length of half diagonal) - (radius of spheres).  
 $\downarrow$  grows as  $\frac{\sqrt{d}}{2}$ .  $\downarrow$  (constant)  $\frac{1}{2}$ .

if  $d > 4$ , inner sphere starts poking outside. ( $\frac{\sqrt{d}-1}{2} > \frac{1}{2}$ ).

## 1.5. Overfitting & the Bias-Variance Tradeoff.

- estimate risk  $\mathbb{E}_{(Y,X) \sim p} [\ell(Y, f(X))]$  by training error (TE).  $\frac{1}{n} \sum \ell(Y_i, f(X_i))$ .  
 $= \mathbb{E}_n [\ell(Y, f(X))]$
- TE always underestimate prediction error.  
 Better estimation: test error.

### - Bias-Variance Tradeoff:

suppose  $Y = f(X) + \epsilon$ ,  $\mathbb{E}(\epsilon) = 0$ ,  $\text{Var}(\epsilon) = \sigma^2$ .

$f(X, T)$ : predicted function trained on  $T$ .

$$\begin{aligned} R(f) &= \mathbb{E}[(Y - f(X, T))^2] \\ &= \mathbb{E}[(Y - f(X))^2] + \mathbb{E}[(f(X) - f(X, T))^2] \\ &= \sigma^2 + \mathbb{E}_X[(f(X) - \mathbb{E}_T[f(X, T) | X])^2] + \mathbb{E}_X[\mathbb{E}_T[(f(X, T) - \mathbb{E}_T[f(X, T) | X])^2 | X]] \\ &= \sigma^2 + \mathbb{E}_X[\text{Bias}_T(X)^2 + \text{Variance}_T(X)]. \end{aligned}$$

$\mathbb{E}[(Y - f(X, T))^2]$   
 $= \mathbb{E}[g(X)(Y - f(X))]$   
 $= \mathbb{E}_X[g(X) \cdot \underbrace{\mathbb{E}_T[Y - f(X) | X]}_{=0}]$   
 $\star$   
 $\mathbb{E}_T[\mathbb{E}_T[(f(X, T) - \mathbb{E}_T[f(X, T) | X])^2 | X]]$

- k-NN)  $f(x) = f(x, T) = \frac{1}{k} \sum_{i=1}^k (f(x_{(i)}) + \epsilon_{(i)})$ . ( $\ell_2$ ):  $L^2$ -N.N.

$$\Rightarrow \text{Bias}_T(x) = f(x) - \frac{1}{k} \sum_{i=1}^k f(x_{(i)}).$$

$$\text{Variance}_T(x) = \frac{1}{k^2} \text{Var}(\sum \epsilon_{(i)}) = \frac{\sigma^2}{k}.$$

$\Rightarrow$  if  $k \uparrow$ , bias  $\uparrow$ , variance  $\downarrow$ .

$$\begin{aligned} \mathbb{E}[(f(x) - f(x, T))^2] \\ &= \mathbb{E}_X[\mathbb{E}_T[(f(x) - f(x, T))^2 | X]] \\ &= \mathbb{E}_X[\underbrace{(f(x) - \mathbb{E}_T[f(x, T) | X])^2}_{\rightarrow \text{bias}} | X] + \mathbb{E}_X[\mathbb{E}_T[(f(x, T) - \mathbb{E}_T[f(x, T) | X])^2 | X]]. \end{aligned}$$

## 2. Linear Classifiers.

### 2.1. Introduction

- decision boundary (between class  $i$  and  $j$ )

$$\{x \in X : \text{for all } r > 0, B(x, r) \cap \mathcal{F}^{-1}\{i\} \neq \emptyset, B(x, r) \cap \mathcal{F}^{-1}\{j\} \neq \emptyset\}$$

$\Rightarrow$   $\mathcal{H}^1$  decision boundary: Union over all  $(i, j)$ .

- linear classifier: decision boundary btw  $i$  and  $j$  given as subset of affine hyperspace.

$$\{x : \underbrace{H(x)} = 0\}, \text{ where } H(x) = \beta_0 + x\beta.$$

- model joint probability dist.

model conditional density function  $P(\text{class} | X)$ .

### 2.2. Bayes Classifiers

$$\text{thm } R(f) = \mathbb{E}_{(Y,X) \sim p} [\ell(Y, f(X))] = \mathbb{E}_X[\mathbb{E}_{Y|X}[\ell(Y, f(X)) | X]].$$

$\rightarrow f^*(x) = \underset{k}{\operatorname{argmax}} P(Y=k | X=x)$ . : Bayes Classifier.

- Binary)  $f^*(x) = \underset{k}{\operatorname{argmax}} (P(Y=1 | X=x) > \frac{1}{2})$ .

$$\begin{aligned} \mathbb{E}[\sum_{j=0}^K \ell(j, f(x)) \mathbb{1}\{f(x)=j\} | X=x] \\ \rightarrow \sum \ell(j, f(x)) \mathbb{E}[\mathbb{1}\{f(x)=j\} | X=x] \\ = \sum \ell(j, f(x)) P(Y=j | X=x). \end{aligned}$$

$$\begin{aligned}
 - \text{ 归约) } P(Y=j | X=x) &= \frac{P(Y=j, X=x)}{P(X=x)} \\
 &= \frac{P(Y=j, X=x)}{\sum P(Y=l, X=x)} = \frac{\pi_j P_j(x)}{\sum \pi_l P_l(x)}.
 \end{aligned}$$

$$\Rightarrow P(Y=k | X=x) > P(Y=j | X=x) \Leftrightarrow \frac{P_k(x)}{P_j(x)} > \frac{\pi_j}{\pi_k}.$$

$$\text{Bayes rule: } f^*(x)=k \text{ when } \frac{P_k(x)}{P_j(x)} > \frac{\pi_j}{\pi_k} \text{ for all } j \neq k.$$

$$\left( f^*(x) = \arg \max \pi_j P_j(x) \right) \rightarrow R(f^*) = \text{Bayes risk}.$$

$$- f^0 \in H \text{ that minimizes } R(f) \Rightarrow \text{oracle classifier: } R(f^0) = \inf_{f \in H} R(f).$$

### 2.3 Logistic Classification.

$$- f(x)=k \text{ if } \pi_k(x) > \pi_j(x) \text{ for all } j \neq k, \quad \pi_j(x) = P(Y=j | X=x) = \mathbb{E}[I(Y=j) | X=x].$$

$$- p(y; \pi) = \pi^y (1-\pi)^{1-y} \text{ for } y=0,1.$$

$$\rightarrow \mathcal{L}(\pi) = \prod_{i=1}^n p(Y_i; \pi) = \prod \pi_i^{Y_i} (1-\pi_i)^{1-Y_i}.$$

$$- \text{Assumption: } P(Y=k|x) = \frac{\exp(\beta_{k0} + x^T \beta_k)}{1 + \sum_{j=1}^{K-1} \exp(\beta_{j0} + x^T \beta_j)}.$$

$$- \text{Bayes classifier: } f^*(x) = \arg \max_j P(Y=j|x).$$

$$- \text{linear boundary: where } P(Y=2|x=x) = P(Y=1|x=x).$$

$$\Rightarrow \exists x \in \mathbb{R}^d: (\beta_{20} - \beta_{10}) + x^T (\beta_2 - \beta_1) = 0.$$

- motivation ①

$$P(Y=1|x) = F(\beta_0 + x^T \beta) \rightarrow F: \text{anti, increasing, } F(-\infty)=0, F(\infty)=1.$$

$$\text{logistic: } F(x) = \frac{\exp(x)}{1 + \exp(x)}.$$

motivation ②

$$\Rightarrow P(Y=2|x=x) = P(Y=1|x=x).$$

$$= \beta_0 + x^T \beta.$$

$$\exists x: \Pr(Y=1|x=x) = 0.5 \Leftrightarrow \exists x: \log \left( \frac{P(Y=1|x=x)}{P(Y=2|x=x)} \right) = 0.$$

$$\text{Suppose log-odds is linear. } \rightarrow P(Y=1|x=x) = \frac{\exp(\beta_0 + x^T \beta)}{1 + \exp(\beta_0 + x^T \beta)}.$$

$$- \text{estimation) mle: } \mathcal{L}(\hat{\beta}_0, \hat{\beta}) = \prod \Pr(Y_i = Y_i | X_i = X_i).$$

$$\text{if binary, } = \prod \pi_i(x_i, \beta)^{Y_i} (1 - \pi_i(x_i, \beta))^{1-Y_i}.$$

$$\begin{aligned}
 \mathcal{L}(\beta_0, \beta) &= \sum \left[ Y_i \log(\pi_i) - (1-Y_i) \log(1-\pi_i) \right] \\
 &= \sum \left[ Y_i (\beta_0 + x_i^T \beta) - \log(1 + \exp(\beta_0 + x_i^T \beta)) \right].
 \end{aligned}$$

$\hookrightarrow$  Newton's method.

$$\frac{\exp(\beta_0 + x_i^T \beta)}{1 + \exp(\beta_0 + x_i^T \beta)}.$$

## 2.4. LDA.

— Let  $p_j(x) = p(x|y=j)$ .  $j=0, \dots, k-1$ .

$$\pi_j = \Pr(y=j) \dots \sum_{j=0}^{k-1} \pi_j = 1.$$

(\* Bayes classifier:  $f^*(x) = \operatorname{argmax}_j p_j(x) \pi_j$ .)

— Model  $p_j(x) = p(x|Y=j)$  as multivariate Gaussian:  $X|Y=j \sim N(\mu_j, \Sigma_j)$ .

Then if  $X|Y=j \sim N(\mu_j, \Sigma_j)$ , then  $\pi_j p_j(x) > \pi_l p_l(x)$

Bayes rule:  $f^*(x) = j$  if  $\pi_j p_j(x) > \pi_l p_l(x)$  for  $\forall l \neq j$ .

$$\text{where } \log \pi_j p_j(x) = -\frac{1}{2} (\log |\Sigma_j| - \frac{1}{2} (x - \mu_j)^T \Sigma_j^{-1} (x - \mu_j) + \log \pi_j).$$

$$\hookrightarrow \log(\pi_j p_j(x))$$

$$\Leftrightarrow f^*(x) = \operatorname{argmax}_j \log(\pi_j p_j(x)).$$

decision boundary:  $\{x \in \mathcal{X} : \log \pi_j p_j(x) = \log \pi_l p_l(x)\} \rightarrow \text{quadratic}.$

— In practice) use sample quantities.

$$\hat{\pi}_j = \frac{n_j}{n}, \quad \hat{\mu}_j = \frac{1}{n_j} \sum_{i: Y_i=j} X_i, \quad \hat{\Sigma}_j = \frac{1}{n_j - 1} \sum_{i: Y_i=j} (X_i - \hat{\mu}_j)(X_i - \hat{\mu}_j)^T.$$

— If we assume  $\Sigma_j = \Sigma$  for all  $j$ .

$$f^*(x) = \operatorname{argmax}_k \log \pi_k p_k(x) \quad \text{where } \log \pi_j p_j(x) = x^T \Sigma^{-1} \mu_j - \frac{1}{2} \mu_j^T \Sigma^{-1} \mu_j + \log \pi_j.$$

decision boundary:  $\{x \in \mathcal{X} : \log \pi_j p_j(x) = \log \pi_l p_l(x)\} \rightarrow \text{linear}.$

$$\hat{\Sigma} = \frac{1}{n-k} \sum_{j=1}^k (n_j - 1) \hat{\Sigma}_j.$$

— LDA:  $d + \frac{d(d-1)}{2}$  parameters vs. QDA:  $d + d(d-1)$ .

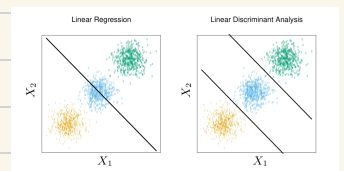
## 2.5. LDA or Logistic Classification.

— LDA need more assumptions.

has troubles with categorical inputs.

useful when some output are missing.

works better for multi-class problems.



### 3. Shrinkage Methods. - Multicollinearity, Shrinkage.

#### 3.1. Regularization.

typically convex.

- $\min_{\beta} \|y - X\beta\|_2^2$  s.t.  $\beta \in C$  (constrained form)
- $\min_{\beta} \|y - X\beta\|_2^2 + P(\beta)$  (penalized form)

- Three norms:  $\| \beta \|_0 = \sum_{j=1}^d \mathbb{I}(\beta_j \neq 0)$
- $\| \beta \|_1 = \sum_{j=1}^d |\beta_j|$
- $\| \beta \|_2^2 = \sum_{j=1}^d \beta_j^2$

- $\min_{\beta} \|y - X\beta\|_2^2$  s.t.  $\| \beta \|_0 \leq k^{(1)} \Leftrightarrow \min_{\beta} \frac{1}{2} \|y - X\beta\|_2^2 + \lambda \| \beta \|_0^{(4)}$
  - " s.t.  $\| \beta \|_1 \leq c^{(2)} \Leftrightarrow \min_{\beta} \frac{1}{2} \|y - X\beta\|_2^2 + \lambda \| \beta \|_1^{(5)}$
  - " s.t.  $\| \beta \|_2 \leq \sqrt{c}^{(3)} \Leftrightarrow \min_{\beta} \frac{1}{2} \|y - X\beta\|_2^2 + \lambda \| \beta \|_2^{(6)}$
- $\lambda \geq 0$  s.t.  $\lambda \rightarrow$

- Toy example)  $Y = (Y_1, \dots, Y_d)^T \in \mathbb{R}^d \sim N_d(\mu, \sigma^2 I)$ .

- MLE of  $\mu = Y$ .

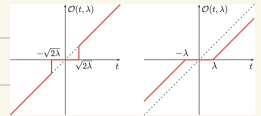
- when  $d \leq 2$ ,  $Y$  is the best estimator.

$$d \geq 3, \text{ best} = \hat{\mu}^{JS} = \left(1 - \frac{(d-2)\sigma^2}{\sum Y_i^2}\right) Y.$$

$$E[\|Y - \mu\|_2^2] = d\sigma^2.$$

$$E[\|\hat{\mu}^{JS} - \mu\|_2^2] = d\sigma^2 - (d-2)^2\sigma^2 E\left[\frac{1}{\|Y\|_2^2}\right].$$

- solutions of  $\hat{\mu}_i = H(Y_i; \sqrt{2\lambda}) \rightarrow$  hard thresholding. dis anti
- $\hat{\mu}_i = S(Y_i; \lambda) \rightarrow$  soft thresholding. conti.
- $\hat{\mu}_i = \frac{Y_i}{1+2\lambda} \rightarrow$  never sparse.



- sparsity: (2):  $\downarrow$  higher degree of sparsity.
- (1), (6):  $\uparrow$

- convexity: (2), (3), (6)  $(p \geq 1)$

#### 3.2. Variable Selection.

- $\hat{\beta}_A = (X_A^T X_A)^{-1} X_A^T y$ ,  $\hat{\beta}_{-A} = 0$ ,  $A = \text{supp}(\hat{\beta})$ .

$$\hat{\beta}_S^{\text{oracle}} = (X_S^T X_S)^{-1} X_S^T y, \hat{\beta}_S = 0, S = \text{supp}(\beta_0), \underline{S_0} = |S|.$$

$$\Rightarrow \text{in-sample risk} : \frac{1}{n} \|X \hat{\beta}^{\text{oracle}} - y\|_2^2 = \sigma^2 \frac{S_0}{n}.$$

risk of oracle risk blz.

$$\frac{1}{n} \mathbb{E} \|X\beta - X\beta_0\|_2^2 \leq \frac{2}{n} (\log d + 2 + o(1)) > \log d.$$

- If  $1 \leq n \log d$ , best subset selection estimator satisfies

$$\begin{aligned} \frac{\mathbb{E} \|X\beta - X\beta_0\|_2^2 / n}{n^2 s / n} &\leq 4 \log d + 2 + o(1) \text{ as } n/d \rightarrow \infty. \\ \inf_{\beta} \sup_{X/\beta_0} \frac{\mathbb{E} \|X\beta - X\beta_0\|_2^2 / n}{n^2 s / n} &\geq 2 \log d - o(\log d). \end{aligned}$$

### 3.3 Ridge Regression.

$$\begin{aligned} \text{Def } \beta_{\text{ridge}} &= \arg \min_{\beta} \sum_{i=1}^n \left( y_i - \beta_0 - \sum_{k=1}^d x_{ik} \beta_k \right)^2 \text{ subject to } \sum_{k=1}^d \beta_k^2 \leq S. \\ &= \arg \min_{\beta} \left\| Y - X\beta \right\|_2^2 + \lambda \|\beta\|_2^2, \quad \beta = (X^T X + \lambda I)^{-1} X^T Y. \end{aligned}$$

### 3.4 LASSO

$$\begin{aligned} \text{Def } \beta_{\text{lasso}} &= \arg \min_{\beta} \sum_{i=1}^n \left( y_i - \beta_0 - \sum_{k=1}^d x_{ik} \beta_k \right)^2 \text{ subject to } \sum |\beta_k| \leq S. \\ &= \arg \min_{\beta} \left\| Y - X\beta \right\|_2^2 + \lambda \sum |\beta_k| \end{aligned}$$

non-differentiable around 0.

- optimization w/ subgradient optimality (KKT conditions).

:  $\beta$  must satisfy  $X^T (y - X\beta) = \lambda S$ . (i)

$$\begin{aligned} \text{where } S \in \partial \|\beta\|_1 &= \begin{cases} \beta_j > 0 & \lambda_j = -1 \\ \beta_j < 0 & \lambda_j = 1 \\ \beta_j = 0 & \lambda_j \in [-1, 1] \end{cases} \quad j = 1, \dots, d. \quad (ii) \end{aligned}$$

→ (i) 의 결과:  $\beta$  가 유일하게 존재하지 않아요,  $S$  는  $X\beta$  의 값이 어느 범위.

→ (ii) 결과,  $S$  의 역할  $\Rightarrow$  overlapped support 일까? 같은 부호.

( $\beta_j > 0$  but  $\beta_j < 0$  for some  $j$  불가).

→  $X$  의  $\text{col}$  general position ( $\approx$  L.I.) 이면  $\beta_{\text{lasso}}$  는 유일?

pf) Let  $A = \text{supp}(\beta)$ .  $S_A = \text{sign}(\beta_A)$ .

$$(i), (ii) \Rightarrow X_A^T (y - X_A \beta_A) = \lambda S_A.$$

$$\Leftrightarrow \beta_A = (X_A^T X_A)^{-1} (X_A^T y - \lambda S_A). \quad \beta_{-A} = 0.$$

$$= (X_A^T X_A)^{-1} X_A^T y - \lambda (X_A^T X_A)^{-1} S_A \quad \left( \begin{array}{l} \text{d. Best subset} \\ \text{shrinkage.} \end{array} \right. \text{ : no shrinkage.}$$

$\Rightarrow$  orthogonal 일 때  $X_A^T X_A = I$ .

$$\beta_j^{\text{lasso}} = \beta_j^{\text{ols}} - \lambda \frac{S_j}{|\beta_j|} \quad \text{if } |\beta_j| > \lambda$$

$$\text{sign}(\beta_j)$$

| 교재 표현            | 의미                                                               |
|------------------|------------------------------------------------------------------|
| Prediction error | $\ X\hat{\beta} - X\beta_0\ ^2$ (확률변수)                           |
| In-sample risk   | $\frac{1}{n} \mathbb{E} \ X\hat{\beta} - X\beta_0\ ^2$           |
| Predictive risk  | $\mathbb{E} (x_0^T \hat{\beta} - x_0^T \beta)^2$ (out-of-sample) |

## 3b. Theoretical Analysis of LASSO.

### 3b.1. Slow Rates.

- **정제 결과**:  $X$ 에 대한 **정제** 값이  $\text{prediction error} + \sqrt{\log d/n}$  정도 **적당**.

- 가정 (2)  $s = \|\beta_0\|_1$ .

$$\cdot \|y - X\hat{\beta}\|_2^2 \leq \|y - X\beta_0\|_2^2 = \|\varepsilon\|_2^2 \quad (\because s = \|\beta_0\|_1 \text{ 이므로 } \beta_0 \text{는 정제 값}).$$

$$\cdot \|y - X\hat{\beta}\|_2^2 = \|X\beta_0 + \varepsilon - X\hat{\beta}\|_2^2 \\ = \|X(\beta_0 - \hat{\beta})\|_2^2 + \|\varepsilon\|_2^2 - 2\langle \varepsilon, X(\beta_0 - \hat{\beta}) \rangle.$$

$$\Rightarrow \|X\hat{\beta} - X\beta_0\|_2^2 \leq 2\langle \varepsilon, X\hat{\beta} - X\beta_0 \rangle = 2\langle X^T \varepsilon, \hat{\beta} - \beta_0 \rangle \\ \leq 4\|\beta_0\|_1 \|\varepsilon\|_\infty \quad (\because \text{Holder's inequality}).$$

By maximal ineq. for Gaussians, for  $\forall s > 0$ ,

$$\max_{j=1, \dots, d} |X_j^T \varepsilon| \leq \sqrt{2n \log\left(\frac{2d}{s}\right)} \quad \text{w.p. at least } 1-s.$$

$$\Rightarrow \frac{1}{n} \|X\hat{\beta} - X\beta_0\|_2^2 \leq 4\|\beta_0\|_1 \sqrt{2 \log\left(\frac{2d}{s}\right)/n} \quad \text{w.p. } " \quad (\Rightarrow \text{Thm}).$$

$$\Leftrightarrow P\left(\frac{1}{n} \|X\hat{\beta} - X\beta_0\|_2^2 \geq \bigcirc\right) \leq s.$$

$$\Rightarrow \mathbb{E}\left[\frac{1}{n} \|X\hat{\beta} - X\beta_0\|_2^2\right] \leq \|\beta_0\|_1 \sqrt{\frac{\log d}{n}} \quad \downarrow \text{ * 이는: lasso-slow rate-proof.}$$

$\underbrace{\sqrt{\log d/n}}_{\text{so } \log d/n.} \quad (\text{cf. subset : } s_0 \log d/n.)$   
 $\hookrightarrow \text{slower.}$

**Lemma.** Suppose we have random variables  $X_1, \dots, X_d$  with  $X_j \sim N(0, \sigma^2)$ , but  $X_j$ 's are not necessarily independent. Then

$$P\left(\max_{j=1, \dots, d} |X_j| \geq t\right) \leq 2d \exp\left(-\frac{t^2}{2\sigma^2}\right).$$

or equivalently, with probability  $1 - \delta$ ,

$$\max_{j=1, \dots, d} |X_j| \leq \sigma \sqrt{2 \log \frac{2d}{1-\delta}}$$

**Theorem.** Suppose the linear model  $Y = X\beta_0 + \varepsilon$  with  $\varepsilon_i$  i.i.d. from  $N(0, \sigma^2)$ . If we choose  $s = \|\beta_0\|_1$  as the tuning parameter, then

$$\frac{1}{n} \|X\hat{\beta} - X\beta_0\|_2^2 \lesssim \frac{\|\beta_0\|_1 \|X^T \varepsilon\|_\infty}{n}.$$

Suppose  $\|X_j\|_2 = \sqrt{n}$ . Then with high probability,

$$\frac{1}{n} \|X\hat{\beta} - X\beta_0\|_2^2 \lesssim \|\beta_0\|_1 \sqrt{\frac{\log d}{n}}.$$

- Out-of-sample risk:  $\mathbb{E}(x_0^T \hat{\beta} - x_0^T \beta)^2 \lesssim \underbrace{\|\beta_0\|_1^2}_{\hookrightarrow \|\beta_0\|^2 \text{ 다 이 쥔}} \sqrt{\frac{\log d}{n}}.$

- Oracle inequality: (X assume linearity of the mean:  $y = \mu(X) + \varepsilon$ ).

• It holds that  $\langle X^T(y - X\hat{\beta}), \tilde{\beta} - \hat{\beta} \rangle, \tilde{\beta} - \hat{\beta} \rangle \leq 0$ . ( $\tilde{\beta}$ : other feasible lasso).

$$\Leftrightarrow \langle \mu(X) - X\hat{\beta}, X\tilde{\beta} - X\hat{\beta} \rangle \leq \langle X^T \varepsilon, \tilde{\beta} - \hat{\beta} \rangle.$$

• Using polarization identity,

$$\|X\hat{\beta} - \mu(X)\|_2^2 + \|X\tilde{\beta} - X\hat{\beta}\|_2^2 \leq \|X\tilde{\beta} - \mu(X)\|_2^2 + 2\langle X^T \varepsilon, \tilde{\beta} - \hat{\beta} \rangle$$

$$\Leftrightarrow \frac{1}{n} \|X\hat{\beta} - \mu(X)\|_2^2 + \frac{1}{n} \|X\tilde{\beta} - X\hat{\beta}\|_2^2 \leq \frac{1}{n} \|X\tilde{\beta} - \mu(X)\|_2^2 + \text{tot} \sqrt{\frac{2 \log(2d)}{n}}$$

$$\Rightarrow \frac{1}{n} \|X\hat{\beta} - \mu(X)\|_2^2 \leq \inf_{\|\tilde{\beta}\|_1 \leq s} \left[ \frac{1}{n} \|X\tilde{\beta} - \mu(X)\|_2^2 \right] + \text{tot} \sqrt{\frac{2 \log(2d)}{n}} \quad \text{w.p. at least } 1-s.$$

A.B.이 inf  $\rightarrow \mu(X)$ 가 linear best 값 0. ( $\mu(X) = X\hat{\beta}^{\text{best}}$ ).

### 3.4.2. Fast Rates.

Under very strong assumptions  $\rightarrow$  faster rates.

$X$  satisfies the restricted eigenvalue condition with const.  $\phi > 0$ .

$$\left[ \text{if } \frac{1}{n} \|Xv\|_2^2 \geq \phi^2 \|v\|_2^2 \text{ for all subsets } J \subseteq \{1, \dots, d\} \text{ s.t. } |J| = s_0 \right. \\ \left. \text{and all } v \in \mathbb{R}^d \text{ s.t. } \|v_{J^c}\|_1 \leq 3\|v_J\|_1 \right]$$

**Theorem.** Suppose  $Y = X\beta_0 + \epsilon$  with  $X$  fixed,  $\epsilon_i$  i.i.d. from  $N(0, \sigma^2)$ , with  $\|\beta_0\|_0 = s_0$ . Suppose  $X$  satisfies the restricted eigenvalue condition in (16). If  $\lambda \geq 2\|X^\top \epsilon\|_\infty$ , then

$$\frac{1}{n} \|X\hat{\beta} - X\beta_0\|_2^2 \lesssim \frac{s_0 \lambda^2}{n^2 \phi_0^2},$$

and

$$\|\hat{\beta} - \beta_0\|_2^2 \lesssim \frac{s_0 \lambda^2}{n^2 \phi_0^2}.$$

Suppose  $\|X_J\|_2 = \sqrt{n}$ , and  $\lambda$  is chosen as  $\lambda = 2\sigma\sqrt{2n \log(2d/\delta)}$ . Then with high probability,

$$\frac{1}{n} \|X\hat{\beta} - X\beta_0\|_2^2 \lesssim \frac{\sigma^2 s_0 \log(d)}{n \phi_0^2},$$

and

$$\|\hat{\beta} - \beta_0\|_2^2 \lesssim \frac{\sigma^2 s_0 \log d}{n \phi_0^2}, \quad (17)$$

with probability tending to 1.

$$\text{pf)} \quad \|y - X\hat{\beta}\|_2^2 + 2\lambda \|\hat{\beta}\|_1 \leq \|y - X\beta_0\|_2^2 + 2\lambda \|\beta_0\|_1. \quad (\because \hat{\beta} \text{ is minimizer})$$

$$\Leftrightarrow \|X\hat{\beta} - X\beta_0\|_2^2 \leq 2\langle \epsilon, X(\hat{\beta} - \beta_0) \rangle + 2\lambda (\|\beta_0\|_1 - \|\hat{\beta}\|_1) \\ \leq 2\|X^\top \epsilon\|_\infty \|\hat{\beta} - \beta_0\|_1 + 2\lambda (\|\beta_0\|_1 - \|\hat{\beta}\|_1) \\ \leq \lambda \|\hat{\beta} - \beta_0\|_1 + 2\lambda (\|\beta_0\|_1 - \|\hat{\beta}\|_1) \quad (\because \text{and})$$

Let  $\hat{\Delta} = \hat{\beta} - \beta_0$  for convenience. Then since  $\beta_{0,-S} = 0$ ,

$$\|\beta_0\|_1 - \|\hat{\beta}\|_1 = \|\beta_{0,S}\|_1 - \|\beta_{0,S} + \hat{\Delta}_S\|_1 - \|\hat{\Delta}_{-S}\|_1 \\ \leq \|\hat{\Delta}_S\|_1 - \|\hat{\Delta}_{-S}\|_1,$$

$$\leq \lambda (\|\hat{\Delta}_S\|_1 + \|\hat{\Delta}_{-S}\|_1) + 2\lambda (\|\beta_0\|_1 - \|\hat{\beta}\|_1) \\ = 3\lambda (\|\hat{\Delta}_S\|_1 + \|\hat{\Delta}_{-S}\|_1) \geq 0.$$

$$\Rightarrow \|\hat{\Delta}_{-S}\|_1 \leq 3\|\hat{\Delta}_S\|_1 \Rightarrow \text{R.E.C. of } \hat{\Delta} = \hat{\beta} - \beta_0 \text{ is satisfied.}$$

$$\|\hat{\Delta}_S\|_1 \leq \sqrt{s_0} \|\hat{\Delta}_S\|_2 \leq \sqrt{s_0} \|\hat{\Delta}\|_2 \quad \text{by } \underbrace{\quad}_{C-S}.$$

$$\|X\hat{\beta} - X\beta_0\|_2^2 \leq 3\lambda \|\hat{\Delta}_S\|_1 \\ \leq 3\lambda \sqrt{s_0} \|\hat{\beta} - \beta_0\|_2. \quad \text{R.E.C.} \\ \leq 3\lambda \sqrt{\frac{s_0}{n\phi_0^2}} \|X\hat{\beta} - X\beta_0\|_2.$$

$$\Leftrightarrow \frac{1}{n} \|X\hat{\beta} - X\beta_0\|_2^2 \leq \frac{9s_0 \lambda^2}{n^2 \phi_0^2} \rightarrow \lambda \geq \|X^\top \epsilon\|_\infty.$$

$$\text{Since } \|X_J\|_2 = \sqrt{n}, \text{ w.p. } 1-\delta, \quad \|X^\top \epsilon\|_\infty \leq \sigma \sqrt{2n \log(2d/\delta)}.$$

$$\text{choose } \lambda = 2\sigma \sqrt{2n \log(2d/\delta)}$$

$$\Rightarrow \frac{1}{n} \|X\hat{\beta} - X\beta_0\|_2^2 \leq \frac{12\sigma^2 \log(2d/\delta)}{n\phi_0^2} \quad \left( \text{order: } \log \frac{d}{n} \right).$$

Coefficients)  $\|\hat{\beta} - \beta\|_2^2 \leq \frac{1}{n\lambda^2} \|X\hat{\beta} - X\beta\|_2^2$

$$\leq \frac{9s \cdot \lambda^2}{n^2 \phi_0^4} = \frac{12 \cdot \lambda^2 s \log(2d/s)}{n \phi_0^4}$$

$\phi_0$ : eigenvalue of  $\Phi$ .  
 $\phi_0$ : smallest eigenvalue of  $\Phi$ .

### 3.4.3. Support Recovery.

Additional Assumptions:

(mutual coherence)  $\|(X_S^T X_S)^{-1} X_S^T X_j\|_1 \leq 1 - \gamma$  for  $j \notin S$

(minimum eigenvalue) for some  $C > 0$ ,  $\lambda_{\min}(\frac{1}{n} X_S^T X_S) \geq C$ .

(minimum signal)  $\beta_{0, \min} = \min |\beta_{0,j}| \geq \frac{1}{\sqrt{C}} \|X_S^T X_S\|_1 + \frac{4\sqrt{C}}{\sqrt{C}}$

$$\Rightarrow P(\text{support}(\hat{\beta}) = \text{support}(\beta)) \rightarrow 1.$$

## 4. Basis Expansion & Kernel Methods.

### 4.1. Introduction.

Basis expansions  $\left[ \begin{array}{l} \text{reg: } f(x) = \sum_{j=1}^q \beta_j h_j(x), \quad h_j: \mathbb{R}^d \rightarrow \mathbb{R}. \\ \text{polynomial reg / regression spline / smoothing spline.} \end{array} \right.$

Kernel smoothing  $\left[ \begin{array}{l} \text{Nadaraya-Watson estimator} \\ \text{local linear reg.} \end{array} \right.$

$\hookrightarrow$   $\hat{f}$  estimator's linear smoother:  $\hat{f}(x) = \sum_i w_i(x) y_i$  (weights  $w_i(x)$  &  $y_i$  at  $x_i$ ).

$\Rightarrow \hat{u} = (\hat{f}(x_1), \dots, \hat{f}(x_n)) = SY$  (for some  $S \in \mathbb{R}^{n \times n}$ ).

$\Rightarrow$  effective df ( $\hat{u}$ ) =  $\text{tr}(S)$ .

$\hookrightarrow$   $h, \lambda$  or not at all.  
 $\lambda$  at  $x_i$

### 4.2. Basis Expansion

$f = \underbrace{h(x)^T}_{\text{basis}} \beta$  :  $\beta \in \mathbb{R}^q$ .  $h: \mathbb{R}^d \rightarrow \mathbb{R}^q$ . (typically  $q \gg d$ ).

#### 4.2.1. Splines.

set of basis functions.

$k$ -th order spline: piecewise polynomial fn of degree  $k$ ,  
 has continuous derivatives of order  $1, \dots, k-1$  at knots  $t_1, \dots, t_p$ .

parameterize splines with  $t_1, \dots, t_p$

$\rightarrow$  truncated power basis  $g_1, \dots, g_{p+k-1}$   $\left\{ \begin{array}{l} g_1(x) = 1, g_2(x) = x, \dots, g_{k+1}(x) = x^k \\ g_{k+j}(x) = (x - t_j)_+^k, \quad j=1, \dots, p. \end{array} \right.$

$(p+1) \times k$  poly -  $(p+k) \times 1$   
 $= (p+1)(k+1) - p k = p k + p + 1$

### 4.2.2. Regression splines

$$f(x) = \sum_{j=1}^{p(k+1)} \beta_j g_j(x)$$

→ define basis matrix  $G$  by  $G_{ij} = g_j(x_i)$ .

$$\hat{\beta} = \arg\min_{\beta \in \mathbb{R}^{p(k+1)}} \|y - G\beta\|_2^2 = (G^T G)^{-1} G^T y$$

$$\hat{f}(x) = \sum \hat{\beta}_j g_j(x)$$

$$\hat{\mu} = (\hat{f}(x_1), \dots, \hat{f}(x_n)) = G(G^T G)^{-1} G^T y. \quad (\text{linear smoother}).$$

### 4.2.3. Natural Splines

—  $f$ : poly-degree  $k$  on  $[t_1, t_2], \dots, [t_{p-1}, t_p]$

$f$ : "  $\frac{k-1}{2}$  on  $[-\infty, t_1]$  and  $[t_p, \infty)$

$f$  is anti & has anti-derivatives of orders  $1, \dots, k-1$  at  $t_1, \dots, t_p$ .

$$\dim(\text{span of } k^{\text{th}} \text{ order N.S.}) = (k+1)(p-1) + \left(\frac{k-1}{2} + 1\right) \cdot 2 - pk = p.$$

### 4.2.4. Smoothing splines

— regularized reg. over natural spline basis. (knots at all inputs  $x_1, \dots, x_n$ ).

motivation)  $x_1, \dots, x_n \in [0, 1]$ . smoothing spline estimate of  $f$

$$\hat{f} = \arg\min_f \sum (y - f(x_i))^2 + \underbrace{\int_0^1 (f^{(m)}(x))^2 dx}_{\substack{\text{2nd order penalty} \\ \text{of } f}} \quad (m = \frac{k+1}{2}).$$

⇒ unique, natural  $k^{\text{th}}$  order spline w/ knots at  $x_1, \dots, x_n$ .

pt) basis:  $n$ -dimensional.  $\rightarrow z_i = f(x_i)$ .  $i=1, \dots, n$ .

We can always find natural spline  $\hat{f}$  w/ knots at  $x_1, \dots, x_n$  that satisfies  $\hat{f}(x_i) = z_i$ .

Define  $h(x) = \hat{f}(x) - f(x)$ .

$$\begin{aligned} \int_0^1 f''(x) h'(x) dx &= f''(x) h'(x) \Big|_0^1 - \int_0^1 f'''(x) h(x) dx && \begin{aligned} f''(0) = f''(1) = 0. \\ f'''(x) = 0 \text{ for } x < x_1 \text{ \& } x > x_n. \end{aligned} \\ &= - \int_{x_1}^{x_n} f'''(x) h(x) dx \\ &= - \sum_{j=1}^{n-1} f'''(x_j) h(x) \Big|_{x_j}^{x_{j+1}} + \int_{x_1}^{x_n} f^{(4)}(x) h(x) dx. && \begin{aligned} f^{(4)}(x) = 0. \\ \rightarrow f^{(4)}(x) = 0. \end{aligned} \\ &= - \sum_{j=1}^{n-1} f'''(x_j) (h(x_{j+1}) - h(x_j)). && \rightarrow h(x_j) = 0 \text{ for all } j=1, \dots, n. \\ &= 0 \end{aligned}$$

$$\begin{aligned} \int_0^1 \tilde{f}''(x)^2 dx &= \int_0^1 (f''(x) + h''(x))^2 dx \\ &= \int_0^1 f''(x)^2 dx + \int_0^1 h''(x)^2 dx + 2 \int_0^1 f''(x) h''(x) dx. \\ &= \int_0^1 f''(x)^2 dx + \int_0^1 h''(x)^2 dx. \\ &\geq \int_0^1 f''(x)^2 dx. \quad (= \Leftrightarrow \underline{h''(x) = 0 \text{ for all } x \in [0, 1]}.) \\ &\Leftrightarrow \underline{h \text{ linear.}} \quad \exists \quad h = 0. \end{aligned}$$

## 4.2.b. Finite Dimensional Form

— Reparameterization) choosing basis  $\eta_1, \dots, \eta_n$

$$\Rightarrow \hat{\beta} = \arg\min_{\beta \in \mathbb{R}^n} \left( y_2 - \sum_{j=1}^n \beta_j \eta_j(x) \right)^2 + \lambda \int_0^1 \left( \sum_{j=1}^n \beta_j \eta_j^{(m)}(x) \right)^2 dx.$$

( $\rightarrow$  finite-dimensional form.)

—  $N$  (basis matrix),  $\Omega$  (penalty matrix)  $\in \mathbb{R}^{n \times n}$ ,

$$\begin{cases} N_{ij} = \eta_j(d_i) \\ \Omega_{ij} = \int_0^1 \eta_i^{(m)}(x) \eta_j^{(m)}(x) dx \end{cases}$$

$$\Rightarrow \hat{\beta} = \arg\min_{\beta \in \mathbb{R}^n} \|y - N\beta\|_2^2 + \lambda \beta^T \Omega \beta.$$

$$= (N^T N + \lambda \Omega)^{-1} N^T y.$$

$$\hat{\mu} = N(N^T N + \lambda \Omega)^{-1} N^T y \quad \text{linear.}$$

## 4.2.b. Pénché Form

$$- (ii) = N(N^T N + \lambda \Omega)^{-1} N^T y$$

$$= N(N^T(I + \lambda \underbrace{N^T \Omega N}_{Q \text{ (p.s.d.)}})N)^{-1} N^T y$$

$$= (I + \lambda Q)^{-1} y$$

— S of eigen decomposition:  $S = \sum_{j=1}^n \frac{1}{1 + \lambda d_j} u_j u_j^T$  ( $Q = U D U^T$ ,  $D = \text{diag}(d_1, \dots, d_n)$ )

$\rightarrow$  orth. basis  $u_1, \dots, u_n \in \mathbb{R}^n$  s.t.  $\sum_{j=1}^n u_j u_j^T = I$ .

shrinkage:  $\exists$  eigenvalue of  $Q$   $\rightarrow$  eigenvector shrink.

## 4.3. Kernel Smoothing

— Kernel fun  $K: \mathbb{R}^d \rightarrow \mathbb{R}$ .

$$\begin{cases} \int K(t) dt = 1 \\ \int t K(t) dt = 0 \\ 0 < \int t^2 K(t) dt < \infty. \end{cases}$$

— Nadaraya-Watson estimate (각 점에 대해  $\hat{f}_h$   $\frac{1}{n}$  locally fit).

$$\hat{f}_h(x) = \frac{\sum_{i=1}^n K\left(\frac{x - x_i}{h}\right) y_i}{\sum_{i=1}^n K\left(\frac{x - x_i}{h}\right)} = \sum w_i(x) y_i = \arg\min_{\theta} \sum_{i=1}^n K\left(\frac{x - x_i}{h}\right) (y_i - \theta)^2$$

## 4.3.1. Error Analysis

— Kernel smoothing estimator: universally consistent ( $\mathbb{E} \|\hat{f}_n - f\|_2^2 \rightarrow 0$  as  $n \rightarrow \infty$ ).

+)  $\hat{f}_n$   $\hat{f}_n$

**Theorem.** Suppose that  $d = 1$  and that  $f_0''$  is bounded. Also suppose that  $X$  has a non-zero, differentiable density  $p$  and that the support is unbounded. Then, the risk is

$$R_n = \frac{h_n^4}{4} \left( \int x^2 K(x) dx \right)^2 \int \left( f_0''(x) + 2f_0'(x) \frac{p'(x)}{p(x)} \right)^2 \frac{1}{p(x)} dx + \frac{\sigma^2 \int K^2(x) dx}{nh_n} \int \frac{dx}{p(x)} + o\left(\frac{1}{nh_n}\right) + o(h_n^4) = o\left(h_n^4 + \frac{1}{nh_n}\right).$$

where  $p$  is the density of  $P_X$ .

$$\rightarrow h_n \text{ of } \text{rate} \quad h_n^4 \uparrow \leftrightarrow \frac{1}{n h_n} \downarrow. \Rightarrow h_n^4 = \frac{1}{n h_n}, \quad h_n = n^{-\frac{1}{5}}.$$

$$\Rightarrow R_n = O(n^{-\frac{4}{5}}). \quad [d\text{-th order: } O(n^{-\frac{d}{d+1}})]$$

$$\rightarrow \text{support of boundary kernel bias order } O(h) \text{ near boundary.} \Rightarrow R_n = O(h^3).$$

$(O(h^2) \text{ interior})$

$\Leftrightarrow$  minimax rate over  $H_d(d, L)$

$$\inf_{\hat{f}} \sup_{f \in H_d(d, L)} \mathbb{E}_P [\int (f(x) - \hat{f}(x))^2] \gtrsim n^{-\frac{2d}{2d+1}}.$$

$\hookrightarrow$  ~~이런 양을 나타내는~~ kernel reg. best possible.

## 4.9.2. Local Linear Regression.

— local constant fit  $\rightarrow$  local linear fit.

at each input  $x$ , define  $\hat{f}(x) = \hat{\theta}_x \leftarrow \arg \min_{\theta} \sum_{i=1}^n K\left(\frac{x - x_i}{h}\right) (y_i - \theta)^2$

$\rightarrow$  consider  $\hat{f}(x) = \hat{\alpha}_x + \frac{1}{p_x} x \leftarrow \arg \min_{\alpha, \beta} \sum_{i=1}^n K\left(\frac{x - x_i}{h}\right) (y_i - \alpha - \beta x_i)^2.$

$\Leftrightarrow (x - \beta p)^T \Omega (x - \beta p)$

$\Rightarrow \hat{f}(x) = \underline{b(x)^T (B^T \Omega B)^{-1} B^T \Omega y}$  where  $b(x) = (1, x) \in \mathbb{R}^{d+1}.$

$w(x) = \underline{\Omega B (B^T \Omega B)^{-1} b(x)}$   
 $\Rightarrow$  linear.

$B \in \mathbb{R}^{n \times (d+1)}$  with  $i$ th row  $b(x_i).$   
 $\Omega$ : diag,  $\Omega_{ii} = K\left(\frac{x - x_i}{h}\right).$

$\hat{y} = (w(x_1)^T y, \dots, w(x_n)^T y) = B (B^T \Omega B)^{-1} B^T \Omega y = S y.$

— local fit reduces bias?

$\mathbb{E}[\hat{f}(x)] = \sum_{i=1}^n w_i(x) f_0(x_i).$   $\hookrightarrow$  Taylor expansion of  $f_0$  about  $x$

$= f_0(x) \sum_{i=1}^n w_i(x) + \nabla f_0(x)^T \sum_{i=1}^n (x_i - x) w_i(x) + R.$

$\stackrel{(*)}{=} f_0(x) + R.$

$\Rightarrow \hat{f}^1$  is unbiased to first-order. (bias  $\neq$  2nd order term of  $\phi_2$ ).

— Bias: universal kernel  $h^2$  of order. (bias  $\neq$   $h^1$ ).

## 4.9.3. Higher-Order Smoothness.

—  $\hat{f}(x) = \hat{\beta}_{x,0} + \sum_{j=1}^k \hat{\beta}_{x,j} x^j \leftarrow \arg \min_{\beta_0, \dots, \beta_k} \sum_{i=1}^n K\left(\frac{x - x_i}{h}\right) (y_i - \beta_0 - \sum_{j=1}^k \beta_j x_i^j)^2$

$= \underline{b(x)^T (B^T \Omega B)^{-1} B^T \Omega y}.$   
 $\hookrightarrow (1, x, \dots, x^k).$

— Taylor expansion  $\mathbb{E}[\hat{f}(x) - f_0(x)]^2 \leq n^{-\frac{2d}{2d+1}}$  일지 모르겠다.

— higher-order kernel:  $\int K(t) dt = 1, \quad \int t^j K(t) dt = 0, \quad 0 < \int |t|^k K(t) dt < \infty.$   
 $\hookrightarrow$  ~~이런~~ kernel reg. + higher-order kernel.  
 $j=1, \dots, k-1.$

# Smoothing Splines & Local Polynomial.

[2024 mid #3]  $f_0$ : cubic smoothing spline

(a) origin  $\rightarrow f_0, f_{CSS}$  mid.

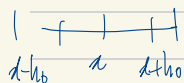
(b)  $E[\hat{f}^p(x)] = f_0(x)$ . : local polynomial is zero  $\Rightarrow$  bias = 0.

$$p) E[\hat{f}^p(x)] = E[\sum w_i(x) y_i] \\ = \sum w_i(x) f_0(x_i) \cdot \mathbb{I}(|x - x_i| \leq h).$$

$$[x - h_0, x + h_0] \cap [t_1, \dots, t_m] = \emptyset \text{ if } h_0 > 0 \text{ small.}$$

$t_1$   $t_2$

$$f_0 \text{ Expansion: } f_0(x_i) = \sum_{j=0}^3 \frac{(x_i - x)^j}{j!} f_0^{(j)}(x). \quad (\because f_0^{(4)} = 0).$$



$$E[\hat{f}^p(x)] = \sum w_i(x) \cdot \mathbb{I}(|x - x_i| \leq h).$$

$$= \sum_{j=0}^3 \frac{f_0^{(j)}(x)}{j!} \sum w_i(x) (x_i - x)^j$$

$$\parallel \\ \left( f_0(x) + f_0'(x) + \frac{1}{2} f_0''(x) + \frac{1}{6} f_0'''(x) \right) \\ \parallel \\ 0.$$

$$= f_0(x).$$

[2025 mid #3]  $\hat{f}^{CSS}$

$\forall x, y \in [0, 1].$

$$(a) \hat{f}^{CSS} \in H_1(3, L) \text{ with prob. } 1 \Leftrightarrow |\hat{f}^{CSS''}(x) - \hat{f}^{CSS''}(y)| \leq L|x - y|.$$

$$p) \text{ if } [t_1, t_{i+1}] \text{ on } H_1, \hat{f}^{CSS} \text{ is not.} \Rightarrow \hat{f}^{CSS'''} \text{ not.}$$

$$\Rightarrow |\hat{f}^{CSS'''}(x)| \leq L_2 \text{ for some } x \in (t_1, t_{i+1}).$$

$$\Rightarrow |\hat{f}^{CSS''}(x)| \leq L \quad (L := \max L_2).$$

$$\Rightarrow |\hat{f}''(x) - \hat{f}''(y)| = \int_y^x \hat{f}'''(u) du \leq L|x - y|.$$

$\rightarrow \forall x \leq y$  on mid.

$$|\hat{f}''(x) - \hat{f}''(y)| \leq | \quad | + \quad | \quad | + \quad | \quad | \\ \leq L|y - x|$$

$$\begin{aligned}
 \text{(b) Bias}(\hat{f}^{\text{IP}}(x)) &= \sum w_i(x) \underbrace{\hat{f}^{\text{GS}}(x_i) - f^{\text{GS}}(x_i)}_{\text{1st Error term}} \\
 &= \sum w_i(x) \left( \frac{1}{j!} \frac{(x_i - x)^j}{j!} f^{(j)}(x) + \frac{(x_i - x)^j}{2!} f^{\text{GS}''}(\xi_i) \right) - f. \\
 &= \sum w_i(x) \frac{(x_i - x)^2}{2} \left( (f^{\text{GS}''})'(\xi_i) - (f^{\text{GS}''})'(x) \right). \\
 &\leq L |\xi_i - x| \leq L |x_i - x|. \\
 &\leq \frac{L}{2!} |w_i(x)| |x_i - x|^2 \mathbb{1}(|x_i - x| \leq h). \\
 &\leq \frac{L}{2!} h^2 \underbrace{|w_i(x)|}_{= C} = C.
 \end{aligned}$$

(c) gel dn of gel.

(d) Bias = 0 on  $[t_2 + h, t_{\text{end}} - h]$ .

$$\begin{aligned}
 \int_0^1 \text{Bias}^2 dx &\leq \int_{\text{total}} \mathbb{1}_{[t_2 + h, t_{\text{end}} - h] \cap [a, b]} (Ch^3)^2 dx \\
 &\leq \sum_{i=1}^m \int_{t_i - h}^{t_i + h} C^2 h^6 dx \\
 &\leq h \cdot C^2 h^6 \cdot m.
 \end{aligned}$$

[HW2 #4]  $\hat{f} : \mathbb{R}^p \rightarrow \mathbb{R} = \sum w_i(x) y_i$ .

(a)  $\sum w_i(x) = 1, \quad \sum w_i(x) (x_i - x)^j = 0$

(b)  $f_0 \in H_1(k+1, L) \rightarrow$  constant bias?

(c)  $\mathbb{E}[\hat{f}(x)] = \sum w_i(x) f_0(x_i)$ .  $H_1(k+1, L) \ni f_0 \Rightarrow k \geq 1 \Rightarrow \mathbb{E}$

$$\begin{aligned}
 \text{Bias} &= \sum_{i=1}^n w_i(x) \left( \sum_{j=0}^{k+1} \frac{(x_i - x)^j}{j!} f^{(j)}(x) + \frac{(x_i - x)^k}{k!} f^{(k)}(\xi_i) \right) - f_0(x) \\
 &= \cancel{f(x)} + ( \quad ) \cdot 0 + \left( \sum w_i(x) \frac{(x_i - x)^k}{k!} \left( f^{(k)}(\xi_i) - \underline{f^{(k)}(x)} \right) \right) \cdot \cancel{f(x)} \\
 &\leq L |\xi_i - x|.
 \end{aligned}$$

$|x_i - x| > h \Rightarrow 0 \text{ ok}$

$$\leq \frac{1}{k!} \sum |w_k(x)| h^k \cdot L_h.$$

$$= \frac{L}{k!} h^{k+1} \sum |w_k(x)|.$$

$$\leq \frac{CL}{k!} h^{k+1} \quad |.$$

<LASSO>.

$\nearrow$  R. eigenvalue.

$$\frac{1}{n} \|X_V\|_2^2 \geq \phi_0^2 \|V\|_2^2 \quad \left( \|V_J\|_1 \leq 3 \|V_J\|_1 \right).$$