

<Statistics>

7. RHS

7.1. Review.

- Goal: find f that minimizes expected prediction error.

$$f^0 = \underset{f \in \mathcal{F}}{\operatorname{argmin}} \mathbb{E}_{(Y,X) \sim p} [L(Y, f(X))].$$

- f^0 is unknown $\rightarrow$ estimated by $\hat{f}$

$$\hat{f} = \underset{f \in \mathcal{F}}{\operatorname{argmin}} \mathbb{E}_n [L(Y, f(X))] = \underset{f \in \mathcal{F}}{\operatorname{argmin}} \frac{1}{n} \sum_{i=1}^n L(Y_i, f(X_i)).$$

7.1.1. Hölder & Sobolev Spaces.

- Hölder class $H_{f,\beta,L}$ on $T \subset \mathbb{R}^d$:

$\hookrightarrow g: T \rightarrow \mathbb{R}^d$ is $\beta-1$ times differentiable

and satisfies $|g^{(k)}(y) - g^{(k)}(x)| \leq L \cdot |y - x|$. ($d=1$).

(bounded derivatives).

- Sobolev class: $\tilde{S}(\beta, L)$.

$\hookrightarrow g: T \rightarrow \mathbb{R}$ s.t. $\int_T (g^{(\beta)}(x))^2 dx \leq L^2$

11.2. Hilbert Space.

- Norm: mapping $\|\cdot\|: V \rightarrow (0, \infty)$

- $$\begin{cases} 1. \|x\| = 0 \iff x = 0. \\ 2. \|ax\| = |a| \|x\| \quad \forall a \in \mathbb{R} \\ 3. \|x+y\| \leq \|x\| + \|y\| \end{cases}$$

- normed space: $\mathbb{R}^d \rightarrow \|x\|_p = \left(\sum_{i=1}^d |x_i|^p \right)^{1/p} \geq \|x\|_1$.

- Inner product: mapping $\langle \cdot, \cdot \rangle: V \times V \rightarrow \mathbb{R}$

- $$\begin{cases} 1. \langle x, x \rangle \geq 0 \text{ and } \langle x, x \rangle = 0 \iff x = 0. \\ 2. \langle ax+by, z \rangle = a \langle x, z \rangle + b \langle y, z \rangle. \\ 3. \langle x, y \rangle = \langle y, x \rangle. \end{cases}$$

- inner product defines a norm: $\|v\| = \sqrt{\langle v, v \rangle}$.

- complete: every Cauchy seq. converges to a limit.

complete,
normed

complete,
inner product

Banach.

$\square$

Hilbert

$\left(\begin{array}{l} f_n \rightarrow f \\ \Rightarrow \|f_n - f\| \rightarrow 0 \text{ as } n \rightarrow \infty \end{array} \right)$ $\uparrow$ iff RHS nil. $\square$

- If V is Hilbert space & L is a closed subspace,
for $v \in V$, $\exists$ unique $y \in L$. (projection of v onto L).
 $\rightarrow y$ minimizes $\|v - z\|$ over all $z \in L$.
 $\rightarrow v$ can be written as $v = w + z$. ($w \in L^\perp$, $z =$ projection of v).

- orthogonal sum: $L \oplus M = \{l + m : l \in L, m \in M\}$.
- every Hilbert space has an orthonormal basis. $\mathcal{E}$.
- Hilbert space is separable if $\exists$ countable orthonormal basis

Thm Let V : separable Hilbert space $\rightarrow \exists$ orthonormal basis $e_1, e_2, \dots$.
 then for any $x \in V$, $x = \sum_{j=1}^{\infty} \theta_j e_j$ where $\theta_j = \langle x, e_j \rangle$.
 and $\|x\|^2 = \sum_{j=1}^{\infty} \theta_j^2$.

Thm (Riesz Representation theorem).

In a Hilbert space X , for any continuous linear functional $L: X \rightarrow \mathbb{R}$,
 $\exists$ unique $y \in X$ s.t. $Lx = \langle x, y \rangle$.

11.4. L_p spaces

- f : functions $[a, b] \rightarrow \mathbb{R}$.
- L_p norm: $\|f\|_p = \left(\int_a^b |f(x)|^p dx \right)^{1/p} \rightarrow \langle f, g \rangle = \int_a^b f(x)g(x) dx$.
- C-S: $\left(\int_a^b f(x)g(x) dx \right)^2 \leq \int_a^b f^2(x) dx \int_a^b g^2(x) dx$.
- Minkowski: $\|f+g\|_p \leq \|f\|_p + \|g\|_p$.
- Hölder: $\|fg\|_1 \leq \|f\|_p \|g\|_q$ ($\frac{1}{p} + \frac{1}{q} = 1$).

- Special properties of $(L_2) \rightarrow$ Hilbert space. $\langle f, g \rangle = \int_a^b f(x)g(x) dx$.
 $\rightarrow$ orthonormal basis $\phi_1, \phi_2, \dots$

$$\begin{aligned} \rightarrow f(x) &= \sum_{j=1}^{\infty} \theta_j \phi_j(x) \\ &\hookrightarrow \theta_j = \int_a^b f(x) \phi_j(x) dx. \\ \int_a^b f^2(x) dx &= \sum_{j=1}^{\infty} \theta_j^2 \end{aligned}$$

$$\rightarrow \text{span of } \{\phi_1, \dots, \phi_n\} = \left\{ \sum \theta_j \phi_j(x) : \theta_1, \dots, \theta_n \in \mathbb{R} \right\}$$

→ projection of $f(x) = \sum_{j=1}^n \alpha_j \phi_j(x)$ onto the span of $\{\phi_1, \dots, \phi_n\}$
 $f_n = \sum_{j=1}^n \alpha_j \phi_j(x)$, (n-term linear approximation of f).

7.4. Sobolev space.

— weakly differentiable if $\exists f'$ st. $\int_x^y f'(s) ds = f(y) - f(x)$.

— Sobolev space of order m

$$S_{m,p} = \{f \in L_p(0,1) : \|D^m f\| \in L_p(0,1)\}. \quad (p=2)$$

Thm S_m is a Hilbert space under the inner product

$$\langle f, g \rangle = \sum_{k=0}^{m-1} f^{(k)}(0) g^{(k)}(0) + \int_0^1 f^{(m)}(x) g^{(m)}(x) dx.$$

→ Define $k(x, y) = \nu$.

Then for each $f \in S_m$, $f(y) = \langle f, k(\cdot, y) \rangle$ and $k(x, y) = \langle k(\cdot, x), k(\cdot, y) \rangle$.

⇒ S_m is a reproducing kernel H.S.

7.6. Mercer kernels and RKHS.

— choose f to minimize $\sum_i (y_i - f(x_i))^2 + \lambda J(f)$.

↔ $f \in \mathcal{M}$ to minimize $\sum (y_i - f(x_i))^2$ where $\mathcal{M} = \{f : J(f) \leq L\}$.

7.6.2. Evaluational functional.

Def $\delta_x : \mathcal{H} \rightarrow \mathbb{R}$ is defined as $\delta_x f = f(x)$.
 ↪ Hilbert space of functions $f: X \rightarrow \mathbb{R}$

Def(1) $\mathcal{H}$ is RKHS if δ_x is continuous for any $x \in X$.

7.6.3. Reproducing Kernel.

Def $K: X \times X \rightarrow \mathbb{R}$ is called a reproducing kernel of $\mathcal{H}$

if $\forall x \in X \quad K_{x\cdot} = K(\cdot, x) \in \mathcal{H}$.

$\forall x \in X$ and $\forall f \in \mathcal{H}, \quad \langle f, K_{x\cdot} \rangle = f(x)$.

→ K is the representer of the evaluational functional.

↳ Recall Thm
 = H.S. of all δ_x
 linear functional is dense.

Def(2) $\mathcal{H}$ is RKHS if it has a reproducing kernel.

7.6.4. Positive semidefinite function.

Def Symmetric function $K: X \times X \rightarrow \mathbb{R}$ is called positive semidefinite

$$\text{iff } \forall n \geq 1, (d_1, \dots, d_n) \in \mathbb{R}^n, (x_1, \dots, x_n) \in X^n, \quad \sum_{i,j=1}^n d_i d_j K(x_i, x_j) = \vec{d}^T K \vec{d} \geq 0.$$

— every inner product
reproducing kernel $\rightarrow$ p.s.d.

— given p.s.d. function K , $K_{xx}(y) = K(y, \underline{x})$

$\rightarrow$ can create fns by linear comb. of the kernel.

$$\rightarrow f(x) = \sum \alpha_j K_{x_j}(x), \quad g(x) = \sum \beta_j K_{y_j}(x)$$

$\rightarrow$ define inner product $\langle f, g \rangle = \langle f, g \rangle_K = \sum \sum \alpha_i \beta_j K(x_i, y_j)$.

$$\text{and norm } \|f\|_K = \sqrt{\langle f, f \rangle_K} = \sqrt{\vec{\alpha}^T K \vec{\alpha}}.$$

Def 2 H_K : completion of H_0 with respect to $\|\cdot\|$. (RHS generated by K).

$\hookrightarrow$ verify: $\langle f, f \rangle \geq 0 \Rightarrow f = 0$.

7.6.5. Mercer Kernels.

Def Mercer Kernel $K(x, y)$: symmetric & p.s.d. $\left(\begin{array}{l} K(x, y) = K(y, x) \\ \int \int K(x, y) f(x) f(y) dx dy \geq 0, \forall f \end{array} \right)$

$\hookrightarrow$ e.g. Sobolev space kernel.

Thm (Mercer's Thm) X : compact, μ : measure on X . $\text{supp}(\mu) = X$.

Suppose $K: X \times X \rightarrow \mathbb{R}$ anti, symmetric, $\sup K(x, y) < \infty$.

Define $T_K f(x) = \int_X K(x, y) f(y) d\mu(y)$

Suppose $T_K f(x)$ is p.s.d. ($= \langle f, T_K f \rangle \geq 0$) $\int_X \int_X K(x, y) f(x) f(y) d\mu(x) d\mu(y) \geq 0$.

$\Rightarrow$ $\exists$ countable eigenvalues & eigenfunctions λ_i, Φ_i , i.e.,

$$\int_X K(x, y) \Phi_i(y) d\mu(y) = \lambda_i \Phi_i(x).$$

$$\begin{aligned} - K(x, y) &= \sum_{i=1}^{\infty} \lambda_i \Phi_i(x) \Phi_i(y) \quad \rightarrow \text{Mercer thm \& Def.} \\ &= \left\langle \sqrt{\lambda_i} \Phi_i(x), \sqrt{\lambda_i} \Phi_i(y) \right\rangle_{\ell^2(\mathbb{N})} \end{aligned}$$

7.6.6. Spectral Representation

— Define feature map Φ

$$\Phi(x) = (\sqrt{\lambda_1} \Phi_1(x), \sqrt{\lambda_2} \Phi_2(x) \dots)$$

$$\Rightarrow K(x, y) = \langle \Phi(x), \Phi(y) \rangle$$

— Expand f in terms of basis $\Phi_1, \Phi_2, \dots$

$$f(x) = \sum_{j=1}^{\infty} \beta_j \Phi_j(x) = \sum_{i=1}^{\infty} d_i K(d_i, x)$$

— If $f(x) = \sum \alpha_j \Phi_j(x)$ and $g(x) = \sum \beta_j \Phi_j(x)$, $\langle f, g \rangle_K = \sum_{j=1}^{\infty} \frac{\alpha_j \beta_j}{\lambda_j}$.

$\hookrightarrow$ since $K(x, \cdot) = K(\cdot, x) = \sum \lambda_i \Phi_i(x) \Phi_i(\cdot)$
 $\langle \Phi_i, \Phi_j \rangle = 0, \quad \because \langle f, K_x \rangle = f(x).$

$$\begin{aligned} \Phi_j(x) &= \langle \Phi_j, K_x \rangle_K = \langle \Phi_j, \sum \lambda_i \Phi_i(x) \Phi_i(\cdot) \rangle \\ &= \lambda_j \Phi_j(x) \underbrace{\langle \Phi_j, \Phi_j \rangle}_1 = \frac{1}{\lambda_j} \end{aligned}$$

7.6.8. Representer Theorem

— let $\hat{f}$ minimize $\ell + \gamma \|f\|_K^2$

then $\hat{f}(x) = \sum_{i=1}^n \alpha_i K(x_i, x)$.

7.6.9. RKHS Regression.

— Define $\hat{f}$ to minimize $R = \sum (y_i - f(x_i))^2 + \gamma \|f\|_K^2$

$\hookrightarrow$ by representer theorem, $\hat{f} = \sum \alpha_i K(x_i, x)$.

$$\rightarrow R = \|Y - K\alpha\|^2 + \gamma \alpha^T K \alpha = \langle Y - K\alpha, Y - K\alpha \rangle + \gamma \alpha^T K \alpha.$$

$\rightarrow$ minimizer: $\alpha = (K + \gamma I)^{-1} Y$.

$$\hat{f}(x) = K\hat{\alpha} = K(K + \gamma I)^{-1} Y = \mathcal{L}Y \rightarrow \text{linear smoother.}$$

7.6.12. Kernel Trick.

— replace $\langle x_i, x_j \rangle$ with $K(x_i, x_j)$

$\Leftrightarrow$ replacing $\langle x_i, x_j \rangle$ with $\langle \Phi(x_i), \Phi(x_j) \rangle = K(x_i, x_j).$

8. Concentration of Measure.

8.2. Notation

$$\begin{aligned} - \mathbb{E} f &= \int f(z) dP(z) = \mathbb{E}[f(Z)] \\ - P_n f &= P_n(f) = \int f(z) dP_n(z) = \frac{1}{n} \sum f(z_i) \end{aligned}$$

8.4. Basic Inequalities

① Hoeffding.

(Markov) $P(Z \geq \varepsilon) \leq \frac{\mathbb{E}(Z)}{\varepsilon}$

(Chebyshev) $P(|Z - \mu| \geq \varepsilon) = P((Z - \mu)^2 \geq \varepsilon^2) \leq \frac{\mathbb{E}(Z - \mu)^2}{\varepsilon^2} = \frac{\sigma^2}{\varepsilon^2}$.

$$\hookrightarrow P(|\bar{Z}_n - \mu| \geq \varepsilon) \leq \frac{\sigma^2}{n\varepsilon^2}$$

$\rightarrow$ Chernoff's method: $P(Z \geq \varepsilon) = P(e^Z \geq e^\varepsilon) = P(e^{tZ} \geq e^{t\varepsilon}) \leq e^{-t\varepsilon} \mathbb{E}(e^{tZ})$

$$P(Z \geq \varepsilon) \leq \inf_{t>0} e^{-t\varepsilon} \mathbb{E}(e^{tZ})$$

Lemma $\mathbb{E}(e^{tZ}) \leq e^{t^2(b-a)^2/8} \quad (a \leq Z \leq b)$.

(Hoeffding) $P(|\bar{Z}_n - \mu| \geq \varepsilon) \leq 2 \exp\left(-\frac{2n\varepsilon^2}{C}\right), \quad (C = n^{-1} \sum (b_i - a_i)^2)$

Corollary w.p. at least $1-\delta$, $|\bar{Z}_n - \mu| \leq \sqrt{\frac{C}{2n \log(\frac{2}{\delta})}} = O\left(\frac{1}{\sqrt{n}}\right)$.

② McDiarmid

(McDiarmid) Suppose $\sup |f(z_1, \dots, z_n) - f(z_1, \dots, z_{i-1}, z_{i+1}, \dots, z_n)| \leq C_i$

Then, $P(|f(z_1, \dots, z_n) - \mathbb{E}(f(z_1, \dots, z_n))| \geq \varepsilon) \leq 2 \exp\left(-\frac{2\varepsilon^2}{\sum C_i^2}\right)$.

$\hookrightarrow$ Thm $P\left(\left|\frac{1}{n} \sum f(z_i) - \mathbb{E}[f(Z)]\right| \geq \frac{\beta}{\sqrt{n}} \left[2 + \sqrt{2 \log\left(\frac{2}{\delta}\right)}\right]\right) \leq \delta$.

$\hookrightarrow$ Ex Let $\bar{F}_n(z) = \frac{1}{n} \sum I(Z_i \leq z)$: empirical cdf.

$$f(z_1, \dots, z_n) = \sup | \bar{F}_n(z) - F(z) | = \frac{1}{n}$$

$\Rightarrow P\left(\sup | \bar{F}_n(z) - F(z) | - \mathbb{E}(\sup | \bar{F}_n(z) - F(z) |) \geq \varepsilon\right) \leq e^{-2n\varepsilon^2}$.

③ Gaussian
Tail Ineq.

Let $X \sim N(0,1)$. $\rightarrow \phi(x) = (2\pi)^{-1/2} e^{-x^2/2}$

$$P(X \geq \varepsilon) = \int_{\varepsilon}^{\infty} \phi(s) ds \leq \frac{1}{\varepsilon} e^{-\varepsilon^2/2}$$

$$\rightarrow P(|\bar{X}_n - \mu| \geq \varepsilon) \leq e^{-n\varepsilon^2/2\sigma^2}$$

$\Rightarrow$ You still need (var. bounded) of X_i .

Bernstein

Lemma $\mathbb{E}(e^{tX}) \leq \exp \left\{ t^2 \sigma^2 \left(\frac{e^c - 1 - tc}{(tc)^2} \right) \right\} \quad \left(|X| \leq c \text{ and } \mathbb{E}(X) = 0 \right).$

$\rightarrow$ (Bernstein) If $P(|X_1| \leq c) = 1$ and $\mathbb{E}(X_1) = 0$,

$$P(|\bar{X}_n| \geq \varepsilon) \leq 2 \exp \left\{ -\frac{n\varepsilon^2}{2\sigma^2 + 2c\varepsilon/3} \right\}.$$

Corollary w.p. at least $1-\delta$, $|\bar{X}_n - \mu| < \sqrt{\frac{2\sigma^2 \log(1/\delta)}{n}} + \frac{2c \log(1/\delta)}{3n}$

2.1. Measures of Complexity $\rightarrow$ uniform bounds $P(\sup |f - \mu(f)| \geq \varepsilon) < \delta$.

- F finite $\rightarrow$ size of the set.

Lemma $P(\max_i Z_i \geq t) \leq \sum P(Z_i \geq t)$. (maximal inequality).

$$\begin{aligned} \rightarrow P(\sup |f(z_1 \dots z_n) - \mu(f)| \geq \varepsilon) &= P(\max |f(z_1 \dots z_n) - \mu(f)| \geq \varepsilon) \\ &\leq \sum P(|f(z_1 \dots z_n) - \mu(f)| \geq \varepsilon). \end{aligned}$$

1) Rademacher Complexity

- $\sigma_1 \dots \sigma_n$: Rademacher uniform variables. $P(\sigma_i = 1) = P(\sigma_i = -1) = \frac{1}{2}$.

Def $\text{Rad}_n(F; \mathbb{Z}^n) = \mathbb{E} \left(\sup \left(\frac{1}{n} \sum \sigma_i f(z_i) \right) \mid z_1, \dots, z_n \right).$

$$\text{Rad}_n(F) = \mathbb{E}(\text{Rad}_n(F; \mathbb{Z}^n)) = \mathbb{E} \left(\sup \left(\frac{1}{n} \sum \sigma_i f(z_i) \right) \right).$$

Lemma Let g : Lipschitz const. $\rightarrow \text{Rad}_n(g \circ F; \mathbb{Z}^n) \leq L \text{Rad}_n(F; \mathbb{Z}^n)$.

2) Shattering Numbers

Def F : class of binary fns. $f \in F: \mathbb{Z} \rightarrow \{0, 1\}$.

$$F_{z_1 \dots z_n} = \{ (f(z_1), \dots, f(z_n)) : f \in F \} : \text{projection of } F \text{ onto } z_1 \dots z_n.$$

Def $s(F, n) = \sup |F_{z_1 \dots z_n}|$

$$\Leftrightarrow s(A, n) : \text{where } F = \{I_A : A \in A\}$$

Thm $\text{Rad}_n(F, \mathbb{Z}^n) \leq \frac{2 \log s(F, n)}{n}$

② VC Dimension

Def $\text{VC}(A)^F = \sup \{n: s(A, n) = 2^n\}$.
 $\hookrightarrow s(F, n) = 2^n$ for $n \leq d \rightarrow \text{poly}$.

Thm suppose F has finite VC dimension d .
 then, $s(F, n) \leq \sum_{i=0}^d \binom{n}{i}$.
 and for $n \geq d$, $s(F, n) \leq \left(\frac{en}{d}\right)^d$.

+) $\exists C$ s.t. $\text{Rad}_n(F) \leq C \sqrt{\frac{1}{n}}$.

§.6. Uniform Bounds

Thm for any $t > \sqrt{\frac{2}{n}}$,

$$P\left(\sup_{f \in F} |(P_n - P)f| > t\right) \leq 4 s(F, 2n) e^{-nt^2/8}$$

$\hookrightarrow$ w.p. at least $1-\delta$, $\sup |P_n(f) - P(f)| \leq \sqrt{\frac{8}{n} \log\left(4 \frac{s(F, 2n)}{\delta}\right)}$.

Lemma (symmetrization) $P\left(\sup |(P_n - P)f| > t\right) \leq 2P\left(\sup |P_n - P_n'|f| > \frac{t}{2}\right)$.

finite VC-dim: $s(F, n) \leq \left(\frac{en}{d}\right)^d$
 $\sup |P_n(f) - P(f)| \leq \sqrt{\frac{8}{n} \left(\log\left(\frac{4}{\delta}\right) + d \log\left(\frac{en}{d}\right)\right)}$.

Lemma $\mathbb{E}\left[\sup |(P_n - P)f|\right] \leq 2 \text{Rad}_n(F)$

Thm w.p. at least $1-\delta$,

$$\begin{aligned} \sup |P_n(f) - P(f)| &\leq 2 \text{Rad}_n(F) + \sqrt{\frac{1}{2n} \log\left(\frac{2}{\delta}\right)} \\ &\leq 2 \text{Rad}_n(F, \mathbb{Z}^n) + \sqrt{\frac{4}{n} \log\left(\frac{2}{\delta}\right)}. \end{aligned}$$

Cor $\sup |P_n(f) - P(f)| \leq \sqrt{\frac{8 \log s(F, n)}{n}} + \sqrt{\frac{1}{2n} \log\left(\frac{2}{\delta}\right)}$.

finite VC-dim $\leq 2C \sqrt{\frac{1}{n}} + \sqrt{\frac{1}{2n} \log\left(\frac{2}{\delta}\right)}$.

8.7. Ex.: Classification

- h_* minimizes $R(h) = \mathbb{P}(Y \neq h(X))$.
- $\hat{h}$ minimizes $\hat{R}_n(h) = \frac{1}{n} \sum \mathbb{I}(Y_i \neq h(X_i))$.

↳ from VC theorem, $\sup |\hat{R}(h) - R(h)| \leq \sqrt{\frac{\log n}{n}}$.

$$\Rightarrow R(h) \leq \hat{R}(\hat{h}) + \sqrt{\frac{\log n}{n}} \leq \hat{R}(h_*) + \sqrt{\frac{\log n}{n}} \leq R(h_*) + \sqrt{\frac{\log n}{n}}.$$

8.8. Ex.: K-means clustering $C = \{c_1, \dots, c_K\}$.
 $C_* = \{c_1^*, \dots, c_K^*\} \rightarrow$ minimizer of $R(C)$.

Def risk $R(C) = \mathbb{E} \|X - \pi_C[X]\|^2$ where $\pi_C[x] = \arg\min \|x - c_j\|^2$.

Thm Suppose $\mathbb{P}(\|X_1\|^2 \leq B) = 1$ for some B .

$$\text{Then, } \mathbb{E}(R(C) - R(C_*)) \leq c \sqrt{\frac{\kappa(d+1) \log n}{n}}.$$

Ⓚ.

9. Minimax

9.1. Definitions

Def Expected risk: $\mathbb{E}_P[d(\theta, \theta(P))]$

Maximum risk: $\sup_{P \in \mathcal{P}} \mathbb{E}_P[d(\theta, \theta(P))]$. $\rightarrow$ ex) def.

Minimax risk: $R_n \approx R_n(P) = \inf_{\theta} \sup_{P \in \mathcal{P}} \mathbb{E}_P[d(\theta, \theta(P))]$
 $\rightarrow$ maximum risk statistic estimator risk.

↳ θ : function or scalar.

Notation $KL(P, P_1) = \int \log \left(\frac{P_0(x)}{P_1(x)} \right) P_0(x) dx$.

— Minimax risk is bound: $L_n \rightarrow$ any estimator of maximum risk.

$L_n \rightarrow L_e \text{ Cam. } \dots$

9.4. Le Cam.

$$\begin{aligned} \text{Then } \textcircled{1} \inf_{\theta} \sup_{P \in \mathcal{P}} \mathbb{E}_P[d(\theta, \theta(P))] &\geq \frac{1}{2} \int [\theta_0^*(x) \wedge \theta_1^*(x)] dx \\ &= \frac{1}{2} [1 - \text{TV}(P_0^*, P_1^*)]. \end{aligned}$$

$$\textcircled{2} \quad \inf_{\theta} \sup_{P \in \mathcal{P}} \mathbb{E}_P [\ell(\hat{\theta}, \theta(P))] \geq \frac{\varepsilon}{8} \exp(-n KL(P_0, P_1)) \\ \geq \frac{\varepsilon}{8} \exp(-n \chi^2(P_0, P_1))$$

$$\textcircled{3} \quad \inf_{\theta} \sup_{P \in \mathcal{P}} \mathbb{E}_P [\ell(\hat{\theta}, \theta(P))] \geq \frac{\varepsilon}{8} \left(1 - \frac{1}{2} \int |P_0 - P_1|\right)^{2n}.$$

Corollary Suppose $\nexists P_0, P_1$ s.t. $KL(P_0, P_1) \leq \log \frac{2}{\alpha}$.

$$\text{then} \quad \inf_{\theta} \sup_{P \in \mathcal{P}} \mathbb{E}_P [\ell(\hat{\theta}, \theta(P))] \geq \frac{\varepsilon}{16}.$$

Ex $(X_1, Y_1) \dots (X_n, Y_n) \rightarrow X_i \sim \text{Unif}(0, 1), Y_i = f(X_i) + \varepsilon_i, \varepsilon_i \sim N(0, 1)$

$$P: P(x, y) = p(x) p(y|x) = p(y - f(x)).$$

$$\text{wlog, } d=0. \Rightarrow \theta = f(0). \quad d(\theta_0, \theta_1) = |\theta_1 - \theta_0|.$$

Let $f_0(x) = 0$ for all x .

$$f_1(x) = \begin{cases} \frac{1}{2}(\varepsilon - x) & 0 \leq x \leq \varepsilon. \\ 0 & x \geq \varepsilon. \end{cases}$$

$$\begin{aligned} \Rightarrow KL(P_0, P_1) &= \int_0^1 \int p_0(x, y) \log\left(\frac{P(x, y)}{P(x, y)}\right) dy dx \\ &= \int_0^1 \int p(y) \log\left(\frac{p(y)}{p(y - f_1(x))}\right) dy dx \\ &= \int_0^\varepsilon KL(N(0, 1), N(f_1(x), 1)) dx. \\ &= \frac{\varepsilon^2}{2} \int_0^1 (1-x)^2 dx = \frac{\varepsilon^2}{6}. \quad \textcircled{X} \end{aligned}$$

5. Ensemble

5.3. Partitions and Trees

- Partitioning range of X : $\rightarrow \{A_1, \dots, A_J\} / x \in A_j$.

$$\hat{h}(x) = \begin{cases} 1 & \text{if } \sum_{x \in A_j} Y_i \geq \sum_{x \in A_j} (1 - Y_i) \\ 0 & \end{cases}$$

$$\hat{f}(x) = \frac{\sum_{j=1}^J \bar{Y}_j \mathbb{I}(x \in A_j)}{\sum_{j=1}^J \bar{Y}_j} \stackrel{\text{o.w.}}{=} \sum \frac{\sum_{i=1}^n Y_i \mathbb{I}(x \in A_j)}{\binom{n}{J}} \mathbb{I}(x \in A_j)$$

$\binom{n}{J} = \sum \mathbb{I}(x \in A_j) \cdot (\text{# of } A_j \text{ that } x \text{ belongs to})$

- Tree of this obj.: impurity $\downarrow$

5.4. Bagging

- Algorithm: $\mathcal{L} = \{ \mathcal{L}^{(b)} : b=1, \dots, B \}$
 $\mathcal{L}^{(b)} : b=1, \dots, B \} \rightarrow f(x, \mathcal{L}^{(b)})$ learned.

- Explanation

$$f_A(x, \mathcal{P}) = \mathbb{E}_{\mathcal{L}} f(x, \mathcal{L}) \rightarrow \text{aggregated predictor.}$$

- avg. pred error in $f(x, \mathcal{L})$ $e = \mathbb{E}_{\mathcal{L}} \mathbb{E}_{X,Y} (Y - f(x, \mathcal{L}))^2$
- avg. " in f_A $e_A = \mathbb{E}_{X,Y} (Y - f_A(x, \mathcal{P}))^2$

$$\Rightarrow \underline{e \geq e_A}$$

- $e - e_A$ depends on $f_A^2(x, \mathcal{P}) = \underbrace{(\mathbb{E}_{\mathcal{L}} f(x, \mathcal{L}))^2}_{e_A} \leq \underbrace{\mathbb{E}_{\mathcal{L}} f(x, \mathcal{L})^2}_e$

$\hookrightarrow f(x, \mathcal{L})$ or $\mathcal{L}$ or both stable $\Rightarrow e \approx e_A$.

$\hookrightarrow f_A$ always improves f (unstable obj. improve $\uparrow$).

- $f_B \leq f_A$ estimator. (if $\mathcal{L} \approx \mathcal{P}$)

$\hookrightarrow$ Bagging unstable learner not $\downarrow$ ≈ 0 .

- Analysis) with $\text{Var}(f(x, \mathcal{L}))$ large?

$$f(x, \mathcal{L}) = \mathbb{I}(\bar{Y}_n \leq x) \text{ and } \text{tree. } (\mathcal{L} = \{Y_1, \dots, Y_n\}).$$

$$\begin{cases} \mathbb{E}[Y_i] = \mu, \text{ Var}[Y_i] = 1. \end{cases}$$

$\left(\begin{array}{l} d \text{ is close to } \mu \text{ relative to } n \\ d = d_n = \mu + \epsilon/\sqrt{n} \end{array} \right.$

$$f(x, z) = \mathbb{P}(\sqrt{n}(\bar{Y}_n - \mu) \leq c) \simeq \mathbb{P}(Z \leq c).$$

$$\rightarrow \text{limiting mean: } \underline{\Phi}(c), \quad \text{variance: } \underline{\Phi}(c)(1 - \underline{\Phi}(c)).$$

$$\begin{aligned} - \text{Bootstrap } \bar{Y}^* &\sim N(\bar{Y}, \hat{\sigma}_Y^2) \\ &\xrightarrow{\text{d}} \sqrt{n}(\bar{Y}^* - \bar{Y}) | Y_1 \dots Y_n \xrightarrow{\text{d}} N(0, 1). \end{aligned}$$

$$\begin{aligned} - f_B(\hat{\mu}_n) &= \mathbb{E}^*[\mathbb{I}(\bar{Y}^* \leq \hat{\mu}_n)] \simeq \mathbb{E}^*[\mathbb{I}(\sqrt{n}(\bar{Y}^* - \bar{Y}) \leq \sqrt{n}(\hat{\mu}_n - \bar{Y}))] \\ &\simeq \underline{\Phi}(\sqrt{n}(\hat{\mu}_n - \bar{Y})) \simeq \underline{\Phi}(cz). \end{aligned}$$

$$\hookrightarrow \frac{\sqrt{n}(\hat{\mu}_n - \mu) + \sqrt{n}(\mu - \bar{Y})}{c}, \quad \frac{N(0, 1)}{\sqrt{n}}.$$

$$\Rightarrow f(x, z) \simeq \mathbb{I}(z \leq c), \quad f_B \simeq \underline{\Phi}(cz), \quad \hookrightarrow 0 \text{ or } 1, \quad \hookrightarrow \text{smoothed ver.}$$

$$(c=0) \quad \rightarrow \text{Ber, mean } \frac{1}{2}, \quad \text{var } \frac{1}{4} \quad \rightarrow \underline{\Phi}(z) = \text{inf}(0, 1), \quad \text{mean } \frac{1}{2}, \quad \text{var } \frac{1}{12}.$$

5.5. Random Forest

5.5.1. Theories.

$$- \text{Suppose each tree } T_k(x) = T(x; \theta_k), \quad \theta_k \stackrel{\text{i.i.d.}}{\sim} \pi(\theta).$$

$$- \text{Let } m_g(x, y) = \frac{1}{n} \sum \mathbb{I}(T_k(x) = y) - \frac{1}{n} \sum \mathbb{I}(T_k(x) = -y). \quad \hookrightarrow y \in \{-1, 1\}.$$

$$PE^*_M = P_{(X, Y)}(m_g(X, Y) < 0).$$

$$PE^*_M \rightarrow PE^* = P_{(X, Y)} \left\{ \mathbb{P}_\theta(T(X) = Y) - \mathbb{P}_\theta(T(X) = -Y) < 0 \right\}$$

SLLN.

We will derive the upper bound of PE^* as a function of the correlation and strength.

Let

$$mr(X, Y) = P_\theta(T(X; \theta) = Y) - P_\theta(T(X; \theta) = -Y)$$

Then the strength of a single tree is

$$s = \mathbb{E}_{(X, Y)}(mr(X, Y)).$$

Chebyshev's inequality gives

$$PE^* \leq \text{Var}(mr)/s^2.$$

Now we calculate $\text{Var}(mr)$. $mr(X, Y)$ can be expanded as

$$\begin{aligned} mr(X, Y) &= \mathbb{E}_\theta(I(T(X; \theta) = Y) - I(T(X; \theta) = -Y)) \\ &= \mathbb{E}_\theta(rmg(\theta, X, Y)), \end{aligned}$$

where $rmg(\theta, X, Y) = I(T(X; \theta) = Y) - I(T(X; \theta) = -Y)$. Note that

$$mr(X, Y)^2 = \mathbb{E}_{\theta, \theta'} rmg(\theta, X, Y) rmg(\theta', X, Y),$$

where θ and θ' are independent and $\theta, \theta' \sim \pi(\theta)$. Then

$$\begin{aligned} \text{Var}(mr) &= \mathbb{E}_{\theta, \theta'} \text{cov}(rmg(\theta, X, Y), rmg(\theta', X, Y)) \\ &= \mathbb{E}_{\theta, \theta'} (\rho(\theta, \theta') sd(\theta) sd(\theta')), \end{aligned}$$

where

$$\rho(\theta, \theta') = \text{corr}_{(X, Y)}(rmg(\theta, X, Y), rmg(\theta', X, Y) | \theta, \theta')$$

and $sd(\theta) = \text{std}_{(X, Y)}(rmg(\theta, X, Y) | \theta)$. Then

$$\text{Var}(mr) = \bar{\rho}(\mathbb{E}_\theta(sd(\theta)))^2 \leq \bar{\rho} \mathbb{E}_\theta(sd^2(\theta))$$

where

$$\bar{\rho} = \mathbb{E}_{\theta, \theta'} (\rho(\theta, \theta') sd(\theta) sd(\theta')) / \mathbb{E}_{\theta, \theta'} (sd(\theta) sd(\theta')).$$

Write

$$\begin{aligned} \mathbb{E}_\theta(sd^2(\theta)) &\leq \mathbb{E}_\theta(\mathbb{E}_{(X, Y)}(rmg(\theta, X, Y)^2)) - \mathbb{E}_\theta(\mathbb{E}_{(X, Y)}(rmg(\theta, X, Y)))^2 \\ &\leq 1 - s^2. \end{aligned}$$

To sum up, we have

$$PE^* \leq \bar{\rho}(1 - s^2)/s^2.$$

It shows that the two ingredients involved in the generalization error for random forests are the strength of the individual classifiers in the forest, and the correlation between them in terms of the raw margin function. It is the first of its kind to explain that an important ingredient in success of ensemble methods is diversity of base learners. For example, random perturbation can be understood as a tool of increasing the diversity.

One example of a simplified tree is the *centered forest* is defined as follows. Suppose the data are on $[0, 1]^d$. Choose a random feature, split in the center. Repeat until there are k leaves. This defines one tree. Now we average M such trees. Breiman (2004) and Biau (2012) proved the following.

Theorem. If each feature is selected with probability $1/d$, $k = o(n)$ and $k \rightarrow \infty$ then

$$\mathbb{E}[f(X) - f^0(X)]^2 \rightarrow 0$$

as $n \rightarrow \infty$.

Under stronger assumptions we can say more:

Theorem. Suppose that f_0 is Lipschitz and that f_0 only depends on a subset S of the features and that the probability of selecting $j \in S$ is $(1/S)(1 + o(1))$. Then

$$\mathbb{E}[f(X) - f^0(X)]^2 = O\left(\frac{1}{n}\right)^{\frac{1}{1 + \frac{1}{2} \frac{1}{\log \frac{1}{S}}}}.$$

This is better than the usual Lipschitz rate $n^{-2/(d+2)}$ if $|S| \leq d/2$. But the condition that we select relevant variables with high probability is very strong and proving that this holds is a research problem.

A significant step forward was made by Scornet, Biau, and Vert (2015). Here is their result.

Theorem. Suppose that $Y = \sum_j f_j(X(j)) + \epsilon$ where $X \sim \text{Uniform}[0, 1]^d$, $\epsilon \sim N(0, \sigma^2)$ and each f_j is continuous, with $f^0(X) = \sum_j f_j(X(j))$. Assume that the split is chosen using the maximum drop in sums of squares. Let t_n be the number of leaves on each tree and let a_n be the subsample size. If $t_n \rightarrow \infty$, $a_n \rightarrow \infty$ and $t_n(\log a_n)^9/a_n \rightarrow 0$ then

$$\mathbb{E}[f(X) - f^0(X)]^2 \rightarrow 0$$

as $n \rightarrow \infty$.

Again, the theorem has strong assumptions but it does allow a greedy split selection. Scornet, Biau, and Vert (2015) provide another interesting result. Suppose that (i) there is a subset S of relevant features, (ii) $p = d$, (iii) m_j is not constant on any interval for $j \in S$. Then with high probability, we always split only on relevant variables.

4.5.4. Variable Importance - $\hat{f}$: random forest est. / X_G data point?

1) LOCO

$$\Delta_j = \frac{1}{m} \sum_{i \in H} W_i \quad (H: \text{hold-out, size } m).$$

$$W_i = (Y_i - \hat{f}_{(-j)}(X_i))^2 - (Y_i - \hat{f}(X_i))^2$$

$\rightarrow \Delta_j$: prediction risk inflation of estimate.

$$\downarrow$$

$$- \mathbb{E}[\Delta_j | \tau] = \mathbb{E}[(Y - \hat{f}_{(-j)}(X))^2 - (Y - \hat{f}(X))^2 | \tau] = \Delta_j.$$

2) Permutation Importance

- permute X_G for the OOB.

• O : oob obs. for tree k .

• O^* : oob " with X_G permuted.

$$- \hat{\Delta}_j = \frac{1}{M} \sum_{k=1}^M W_{kj}$$

$$\text{where } W_{kj} = \frac{1}{m_k} \sum_{i \in O_k^*} (Y_i - \hat{f}_k(X_i))^2 - \frac{1}{m} \sum_{i \in O_k} (Y_i - \hat{f}_k(X_i))^2$$

Estimating

$$\rightarrow \bar{\Delta}_j = \mathbb{E}[(Y - \hat{f}_j(X_j))^2] - \mathbb{E}[(Y - \hat{f}(X))^2].$$

$$\rightarrow \text{Zig-Zag} \quad \bar{\Delta}_j = 2 \left(\frac{1}{1 - \text{obc}_j} \right)^2.$$

4.6. Boosting.

- combining weak learners.

- Let $z_t = (X_t, Y_t)$, $Y_t \in \{-1, +1\}$

- weak learning alg. for some $\gamma > 0$, $\exists$ algorithm that returns $h \in \mathcal{H}$

s.t. for all P , $P(R(h) \leq \frac{1}{2} - \gamma) \geq 1 - \delta$

- Algorithm:

1. Set $P(i) = \frac{1}{n}$. → data point weight.

2. Repeat for $t = 1, \dots, T$

(a) Let $h_t \sim \argmin_{h \in \mathcal{H}} P_t(h(X_i) \neq Y_i)$.

(b) $\epsilon_t = P_t(Y_i \neq h_t(X_i)) = \sum P_t(i) \mathbb{1}(Y_i \neq h_t(X_i))$.

(c) $\alpha_t = \frac{1}{2} \log \frac{1 - \epsilon_t}{\epsilon_t}$.

(d) Let $P_{t+1}(i) = \frac{P_t(i) e^{-\sum_{k=1}^t \alpha_k \mathbb{1}(Y_i \neq h_k(X_i))}}{\sum_j P_t(j) e^{-\sum_{k=1}^t \alpha_k \mathbb{1}(Y_j \neq h_k(X_j))}}$ → $\epsilon_t \rightarrow 0$: error ↓, $\alpha_t \rightarrow \infty$
 $\epsilon_t \rightarrow 1$: error ↑, $\alpha_t \rightarrow -\infty$

3. Set $g(x) = \sum_t \alpha_t h_t(x)$.

4. Return $h(x) = \text{sign } g(x)$.

4.6.2. Training Error.

Lemma
$$Z_t = 2\sqrt{E_t(1-E_t)} = \sum_i D_t(i) e^{-\alpha E_t h_t(x_i)}.$$

Thm Suppose that $\gamma \leq \frac{1}{2} - E_t$ $\forall t$.

Then,
$$R(h) \leq e^{-2\gamma^2 T} \quad (\text{error} \rightarrow 0).$$

4.6.3. Generalization Error.

- $H = \{ \text{sign}(\sum_{t=1}^T h_t(x)) : h_t \in H, h_t \in H \} \rightarrow$ VC dimension T .

$\Rightarrow$ shattering number $\left(\frac{en}{T}\right)^T$.

- If H is finite, $\left(\frac{en}{T}\right)^T \cdot |H|^T$.

infinite & VC dimension b , $\left(\frac{en}{T}\right)^T \left(\frac{en}{2}\right)^{bT} \leq n^{Td}$.

- By the VC thm, w.p. at least $1-\delta$,

$$E[\text{sign}(f(x))] = R(h) \leq \hat{R}(h) + \sqrt{\frac{T \log n}{n}} \rightarrow \text{Too dep. relation!}$$

$\downarrow$ γ is too small.

Thm w.p. at least $1-\delta$,
$$R(g) \leq \hat{R}_\rho(g/\|x\|_1) + \frac{1}{\rho} \sqrt{\frac{2b \log(en/b)}{n}} + \sqrt{\frac{2 \log(2/\delta)}{n}}.$$

Thm
$$\hat{R}(g/\|x\|_1) \leq \prod_{t=1}^T \underbrace{\sqrt{4E_t(1-E_t)}}_{<1}^{HT} \leq b^T.$$

(assuming $\gamma \leq \frac{1}{2} - E_t$).

6. Model Selection and Assessment.

- $\hat{f} \triangleq \underset{f \in \mathcal{F}}{\text{argmin}} E[(Y - f(x))^2] = \underset{f}{\hat{f}} \quad (\neq f \text{ s.t. } Y = f(x)).$
 $\Rightarrow F_1, F_2$ 중 무엇이 더 좋을까?

6.2. Introduction

- expected prediction error : $E[(Y - f(x))^2]$

$\rightarrow$ minimized at $f(x) = \underbrace{E(Y | X=x)}_{\text{true regression function}}.$

- training sample z^n on $\hat{f}$ construct.

- expected prediction (test) error : $E[(Y - \hat{f}(x))^2]$.

$\hookrightarrow$ 1) assessment

- $\left\{ \begin{array}{l} \text{Prediction : cross-validation, AIC, (w)} \\ \text{selecting tune : BIC.} \end{array} \right.$

2) selection.

6.3. Prediction and B-V. Tradeoff.

- expected test error

$$\mathbb{E}[(Y - \hat{f}(x))^2 | X=x] = \sigma^2 + \underbrace{\mathbb{E}[\hat{f}(x) - f(x)]^2}_{\text{Bias}^2(\hat{f}(x))} + \underbrace{\mathbb{E}[(f(x) - \mathbb{E}[\hat{f}(x)])^2]}_{\text{Var}(\hat{f}(x))}.$$

$\mathbb{E}[Y|X]$
training dataset randomness.

→ unbiased OTE is error itself.

6.4. Training Error and Optimization.

- How to estimate $\mathbb{E}P.E?$ $\mathbb{E}[(Y - \hat{f}(x))^2] = R(\hat{f})$.

↳ test data $(x'_1, y'_1), \dots, (x'_n, y'_n)$ $\frac{1}{n} \sum (y'_i - \hat{f}(x'_i))^2$ is not...

- $(x'_1, y'_1), \dots, (x'_n, y'_n) \xrightarrow{\text{i.i.d.}} (x_1, y_1), \dots, (x_n, y_n)$ are i.i.d.,
 $(y_i = f(x_i) + \varepsilon_i, y'_i = \hat{f}(x'_i) + \varepsilon'_i)$.

$$\begin{aligned} \rightarrow \text{Expected Test Error} &= \frac{1}{n} \mathbb{E} \|y' - \hat{f}\|_2^2 \\ &= \frac{1}{n} \mathbb{E} \|y' - f + f - \hat{f}\|_2^2 \\ &= \underbrace{\frac{1}{n} \mathbb{E} \|y - \hat{f}\|_2^2}_{\text{expected training error}} + \underbrace{\frac{2}{n} \text{tr}(\text{Cov}(y, \hat{f}))}_{\text{bias}^2} + \underbrace{\frac{2}{n} \text{tr}(\text{Cov}(y, \hat{f}))}_{\text{optimism}}. \end{aligned}$$

$$\rightarrow \frac{2}{n} \text{tr}(\text{Cov}(y, \hat{f})) = \frac{2}{n} \sum \text{Cov}(y_i, \hat{f}(x_i)).$$

6.5. D.F. and Unbiased Risk Estimation.

- $d_f(\hat{f}) = \frac{1}{n} \sum \text{Cov}(y_i, \hat{f}(x_i)) \xrightarrow{\text{optimism}} = \frac{2\sigma^2}{n} d_f(\hat{f})$.

• # of params. used by the fit $\hat{f}$. (complexity of $\hat{f}$).

e.g. $\hat{f}(x_0) = \bar{y}$, $d_f(\hat{f}) = 1$.

- $\hat{r} = \frac{1}{n} \|y - \hat{f}\|_2^2 + \frac{2\sigma^2}{n} d_f(\hat{f})$: unbiased estimate for the expected test error.

$$\Rightarrow \hat{R} = \hat{f} - \sigma^2 \Rightarrow \mathbb{E}[\hat{R}] = \frac{1}{n} \|f - \hat{f}\|_2^2$$

- (model selection) $\hat{f}_{\hat{\theta}}$ is fitting.

$$\rightarrow \hat{\theta} = \argmin \frac{1}{n} \|y - \hat{f}_{\hat{\theta}}\|_2^2 + \frac{2\sigma^2}{n} d_f(\hat{f}_{\hat{\theta}}).$$

(fit optimistic model has risk?)

6.6. CV

- estimate $E[(y - \hat{f}(x))^2]$ w/ minimal assumptions. ($(x_i, y_i) \sim \text{iid}$)
- split training set into K folds. ($F_1 \dots F_K$)
 - k^{th} fold exclude fit. $\rightarrow \hat{f}^{(-k)}$
 - evaluate errors on the k^{th} fold. $\rightarrow CV_k(\hat{f}^{(-k)}) = \frac{1}{n_k} \sum_{i \in F_k} (y_i - \hat{f}^{(-k)}(x_i))^2$
 - average $\rightarrow CV(\hat{f}) = \frac{1}{K} \sum CV_k(\hat{f}^{(-k)})$. ($K=n$: LOOCV).
 - K of $\hat{f}$: bias-variance trade off.

$$\text{Var}(CV(\hat{f})) = \text{Var}\left(\frac{1}{K} \sum CV_k(\hat{f}^{(-k)})\right) \stackrel{\text{valid for small } K}{\approx} \frac{1}{K} \text{Var}(CV_1(\hat{f}^{(-1)})).$$

$$\rightarrow \frac{1}{K} \sum (CV_k(\hat{f}^{(-k)}) - CV(\hat{f}))^2 \approx \text{estimate} = (SE(\hat{f}_0))^2.$$

6.7. AIC

- KL divergence p_i : true density ; $p_{\hat{\theta}_j} = p(\cdot; \hat{\theta}_j)$: estimated density. $\rightarrow \text{MLE}$.
 $\rightarrow \hat{f} \approx \arg \min_{\hat{f}} KL(p, \hat{f})$
- minimizing $KL(p, p_{\hat{\theta}_j}) \Leftrightarrow$ maximizing $R_j = \int p(x) \log p_{\hat{\theta}_j}(x) dx$.
 $CV(\hat{f}) = \frac{1}{K} \sum \frac{1}{n_k} \sum_{i \in F_k} \log p(x_i; \hat{\theta}_j^{(-k)})$
- $\hat{R}_j = \frac{1}{n} \sum \log p(x_i; \hat{\theta}_j) = \frac{1}{n} l_j(\hat{\theta}_j)$: by likelihood eqs.
 $\downarrow$
 adjust bias: $\hat{R}_j = \frac{1}{n} [l_j(\hat{\theta}_j) - d_j]$.
 $AIC_j = 2l_j(\hat{\theta}_j) - 2d_j$

6.8. BIC

$$BIC = 2l_j(\hat{\theta}_j) - d_j \log n. \quad \left(\frac{n_j}{n} \text{ of true model zero} \rightarrow \text{to smallest contains } p. \rightarrow \hat{f} \rightarrow \hat{f}_j \right)$$

6.9. Linear Smoothers

- $\hat{f}(x) = \sum w_i(x) y_i = w(x)^T y$.
 $\hat{\mu} = S y$. $\rightarrow \hat{f}(\hat{f}) = \frac{1}{n-2} \text{tr}(\text{Cov}(y, \hat{y})) = \text{tr}(S)$.
- Cross Validation: (linear smoother w/ $\hat{f}$, $\hat{f} \rightarrow \hat{g}$)
 $CV(\hat{f}) = \frac{1}{n} \sum (y_i - \frac{y_i - \hat{f}(x_i)}{1 - S_{ii}})^2 = \frac{1}{n} \sum \left(\frac{y_i - \hat{f}(x_i)}{1 - S_{ii}} \right)^2$.
 $\frac{1}{1 - S_{ii}} (\hat{f}(x_i) - S_{ii} y_i)$.