

Chapter IV (alternative)

①
Eddie Hamon
04/10/24

Learning & Minimax theory in a nutshell

I Empirical risk minimization : A simple oracle bound

II Minimax lower bounds : A bayesian method applied to regression

I Empirical risk minimization

We observe $(X_1, Y_1), \dots, (X_n, Y_n)$ iid with distribution $P_{(X,Y)}$ over $\mathcal{X} \times \mathcal{Y}$.

The goal is to learn a function $f: \mathcal{X} \rightarrow \mathcal{Y}$ with small risk

$$\mathcal{R}(f) = \mathbb{E}_P[l(Y, f(X))]$$

where $l: \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}$ is a prescribed loss function

Ex: Classification

$$\mathcal{Y} = \{-1, 1\}$$

$$l(y, y') = \mathbb{1}_{y \neq y'}$$

Regression

$$\mathcal{Y} = \mathbb{R}$$

$$l(y, y') = (y - y')^2$$

(2)

We choose to search f in a prescribed set $\mathcal{H}$ of (measurable) functions

Ex: $\mathcal{H} :=$ linear functions on $\mathcal{X} = \mathbb{R}^d$
:= A class of neural nets
:= sum of N Fourier functions

We try to minimize R by the empirical risk

$$\hat{R}(f) := \frac{1}{n} \sum_{i=1}^n \ell(y_i, f(x_i))$$

The ERM on $\mathcal{H}$ is then

$$\hat{f} \in \underset{f \in \mathcal{H}}{\operatorname{argmin}} \hat{R}(f)$$

The goal is to compare the (expected) performances of $\hat{f}$, compared to the best possible over the class $\mathcal{H}$

$$f^* \in \underset{f \in \mathcal{H}}{\operatorname{argmin}} R(f)$$

We have, almost surely,

$$\underbrace{R(\hat{f})}_{\text{(Random!) (performance)}} = \underbrace{R(f^*)}_{\text{(Deterministic) approximation error}} + \underbrace{R(\hat{f}) - R(f^*)}_{\text{Estimation error}}$$

The approximation error depends on $\mathcal{H}$ and $P_{X,Y}$. In regression, it has been studied in chapter 2. We focus on the estimation term

$$\begin{aligned} R(\hat{f}) - R(f^*) &= (R(\hat{f}) - \hat{R}(\hat{f})) + (\hat{R}(\hat{f}) - \hat{R}(f^*)) \\ &\leq 0 \text{ by definition of } \hat{f} \\ &\quad + (\hat{R}(f^*) - R(f^*)) \\ &\leq \delta \text{ if risk minimization is done approximately} \\ &\leq 2 \sup_{f \in \mathcal{H}} |\hat{R}(f) - R(f)| \end{aligned}$$

From the central limit theorem we know that for fixed $f \in \mathcal{H}$,

$$\sqrt{n} (\hat{R}(f) - R(f)) \xrightarrow[n \rightarrow \infty]{d} N(0, \text{Var}(\ell(X, f(X))))$$

so that $|\hat{R}(f) - R(f)| \leq \frac{1}{\sqrt{n}}$ whp

Here we need a stronger bound holding whp on all $f \in \mathcal{H}$
 $\Rightarrow$ Depends on the complexity of the class $\mathcal{H}$

For simplicity, assume that $Y = \mathbb{R}^p$ (4)

$$\bullet \forall f \in \mathcal{F}, \|f\|_{\infty} \leq M \quad (B)$$

$\bullet l$ is G -Lipschitz in the second variable

$$|l(y, y') - l(y, y'')| \leq G \|y' - y''\| \quad (L)$$

From (L), we get that $\forall f, f' \in \mathcal{F}$,

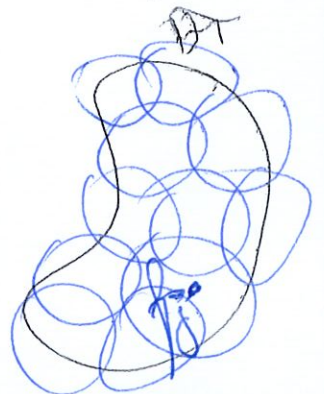
$$\begin{aligned} |\mathcal{R}(f) - \mathcal{R}(f')| &= |\mathbb{E}[l(Y, f(X)) - l(Y, f'(X))]| \\ &\leq G [\mathbb{E} \|f(X) - f'(X)\|] \\ &\leq G \|f - f'\|_{\infty} \end{aligned}$$

and similarly, $|\hat{\mathcal{R}}(f) - \hat{\mathcal{R}}(f')| \leq G \|f - f'\|_{\infty}$

This will allow us to discretize the supremum $\sup_{f \in \mathcal{F}}$ and simplify it to a finite family.

Given $\delta > 0$, write $f_1, \dots, f_N$ for a minimal δ -covering of $\mathcal{F}$ in L^{∞} norm. That is $N = N(\mathcal{F}, \|\cdot\|_{\infty}, \delta)$

$$\mathcal{F} \subset \bigcup_{j=1}^N B_{\|\cdot\|_{\infty}}(f_j, \delta)$$



(5)

For all $f \in \mathcal{F}$, there exists f_j with $\|f - f_j\|_\infty \leq \delta$, and hence

$$\begin{aligned} |R(f) - \hat{R}(f)| &\leq |R(f) - R(f_j)| + |R(f_j) - \hat{R}(f_j)| + |\hat{R}(f_j) - \hat{R}(f)| \\ &\leq 2G \|f - f_j\|_\infty + |R(f_j) - \hat{R}(f_j)| \\ &\leq 2G\delta + |R(f_j) - \hat{R}(f_j)| \end{aligned}$$

Hence,
$$\sup_{f \in \mathcal{F}} |R(f) - \hat{R}(f)| \leq 2G\delta + \max_{j \leq N_\delta} |R(f_j) - \hat{R}(f_j)|$$

Let now f_j be fixed. We assume that the loss function is bounded

$$\boxed{\forall y, y' \in \mathcal{Y}, |l(y, y')| \leq l_\infty}$$

Then
$$R(f_j) - \hat{R}(f_j) = \frac{1}{n} \sum_{i=1}^n (\mathbb{E}[l(Y_i, f_j(X_i))] - l(Y_i, f_j(X_i)))$$

is the sum of iid centered variables bounded by $\frac{l_\infty}{n}$. From Hoeffding's concentration inequality, $\forall t > 0$,

$$\begin{aligned} P(|R(f_j) - \hat{R}(f_j)| > t) &\leq 2 \exp\left(-\frac{2t^2}{n\left(\frac{2l_\infty}{n}\right)^2}\right) \\ &= 2 \exp\left(-\frac{nt^2}{2l_\infty^2}\right) \end{aligned}$$

At the end of the day, we apply a union bound to get

(6)

$$\begin{aligned}
 P\left(\sup_{f \in \mathcal{H}} |R(f) - \hat{R}(f)| > \frac{t}{2} + 2GS\right) &\leq P\left(\prod_{j \leq N_S} |R(f_j) - \hat{R}(f_j)| > t\right) \\
 &\leq N_S \prod_{j \leq N_S} P(|R(f_j) - \hat{R}(f_j)| > t) \\
 &\leq N_S \times 2 \exp\left(-\frac{nt^2}{2L_\infty}\right)
 \end{aligned}$$

En posant $\beta = 2N_S \exp\left(-\frac{nt^2}{2L_\infty}\right)$, on obtient donc qu'avec proba
 au moins $1 - \beta$, $\Leftrightarrow \frac{nt^2}{2L_\infty} = \log\left(\frac{2N_S}{\beta}\right) \Leftrightarrow t = \sqrt{\frac{2L_\infty \log\left(\frac{2N_S}{\beta}\right)}{n}}$

$$\sup_{f \in \mathcal{H}} |R(f) - \hat{R}(f)| \leq 2GS + \sqrt{\frac{2L_\infty \log\left(\frac{2N_S}{\beta}\right)}{n}}$$

Coming back to the initial bound, with proba $\geq 1 - \beta$,

$$\boxed{R(f) \leq R(f^*) + \sqrt{\frac{8L_\infty \log\left(\frac{2N_S}{\beta}\right)}{n}} + 4GS} \quad \text{Oracle bound in probability}$$

Using that $\mathbb{E}[Z] = \int_0^\infty P(Z > t) dt$ if $Z \geq 0$, we get a similar expected bound.

II Minimax lower bound technique

(7)

Given a particular estimator $\hat{\theta}$ that we designed over n -samples $Z_1, \dots, Z_n \sim P$, one may wonder whether one couldn't get a more precise one, for instance on average. Minimax lower bounds aims at doing this, by showing some intrinsic information-theoretic limitations of specific problems

Framework: Model $\mathcal{P}$

- Parameter of interest $\theta: \mathcal{P} \rightarrow \mathcal{H}$
- Metric $d: \mathcal{H} \times \mathcal{H} \rightarrow \mathbb{R}_{\geq 0}$

The minimax risk for estimating $\theta(P)$ over $\mathcal{P}$ is

$$R_n(\mathcal{P}) := \inf_{\hat{\theta}_n} \sup_{P \in \mathcal{P}} E_{\mathcal{P}^n} [d(\theta(P), \hat{\theta}_n)]$$

where $\hat{\theta}_n = \hat{\theta}_n(Z_1, \dots, Z_n)$ ranges among all the estimators over n points $Z_1, \dots, Z_n$

$$\begin{array}{l} \text{Ex: } \mathcal{P} = \{ \mathcal{D}(X, Y) \mid \begin{array}{l} \cdot X \sim \text{Unif}[0, 1]^d \\ \cdot Y = f(X) + \epsilon \\ \cdot X \perp \epsilon \\ \cdot \text{Var}(\epsilon) \leq \sigma^2, E[\epsilon] = 0 \\ \cdot f \in \mathcal{F}^B \end{array} \} \quad \left\{ \begin{array}{l} \cdot \theta(P) = f^* \\ \cdot d(f, f') = \|f - f'\|_{L_2(\mathcal{P}, \mathbb{I}^d)}^2 \end{array} \right. \end{array}$$

Studying $R_n(\mathcal{P})$ is twofold

(8)

Upper bounds: Exhibit as particular $\hat{\theta}_n$ which has good $R_n(\mathcal{P}) \leq \dots$ estimation rate uniformly over $\mathcal{P}$

Lower bounds: Show that one cannot do better, at least up $R_n(\mathcal{P}) \geq \dots$ to constant

For lower bounds, the main idea is to replace the $\sup_{P \in \mathcal{P}}$ by an average $\int_{\mathcal{P}} \pi(dP)$.

Indeed, if π is a probability distribution over $\mathcal{P}$, then for all $\hat{\theta}_n$,
 $\hookrightarrow$ Prior distribution

$$\sup_{P \in \mathcal{P}} \mathbb{E}_{\mathcal{P}^n} [d(\theta(P), \hat{\theta}_n)] \geq \int_{\mathcal{P}} \mathbb{E}_{\mathcal{P}^n} [d(\theta(P), \hat{\theta}_n)] \pi(dP)$$

The right-hand term can usually be studied relatively easily,

since • WE choose π

• The best possible estimator $\hat{\theta}_\pi$ (Bayes estimator) can be explicit

To cover the main ideas, we will present the simplest case where (9)

$$\boxed{\pi = \frac{1}{2}(\delta_{P_0} + \delta_{P_1})}$$

Two hypotheses

Thm: (Le Cam's lemma)

For all $P_0, P_1 \in \mathcal{P}$,

$$R_n(\mathcal{D}) \geq \frac{1}{2} d(\sigma(P_0), \sigma(P_1)) \int P_0^{\otimes n} \wedge P_1^{\otimes n} dV^{\otimes n}$$

where $\|P_0, P_1\| < 1$
 $P_f = \frac{dP_f}{dV}$

Proof: $R_n(\mathcal{D}) \geq \inf_{\hat{\sigma}_n} \frac{1}{2} \left(\mathbb{E}_{P_0^{\otimes n}}[d(\sigma(P_0), \hat{\sigma}_n)] + \mathbb{E}_{P_1^{\otimes n}}[d(\sigma(P_1), \hat{\sigma}_n)] \right)$

$$= \inf_{\hat{\sigma}_n} \frac{1}{2} \int \left(d(\sigma(P_0), \hat{\sigma}_n(\gamma_1, \dots, \gamma_n)) P_0^{\otimes n}(\gamma) + d(\sigma(P_1), \hat{\sigma}_n(\gamma_1, \dots, \gamma_n)) P_1^{\otimes n}(\gamma) \right) dV^{\otimes n}(\gamma)$$

$$\geq \frac{1}{2} d(\sigma(P_0), \sigma(P_1)) \int P_0^{\otimes n} \wedge P_1^{\otimes n} dV^{\otimes n} \quad \square$$

Rk: Actually, $\int P_0 \wedge P_1 = 1 - \frac{1}{2} \int |P_0 - P_1|$ total variation

We can further bound the right-hand term

(10)

Lemma: $\int P_0^{\otimes n} \wedge P_1^{\otimes m} \geq \frac{1}{2} \exp(-n KL(P_0, P_1))$,
 where $KL(P_0, P_1) = \begin{cases} \int \log\left(\frac{P_0}{P_1}\right) P_0 & \text{if } P_0 \ll P_1 \\ \infty & \text{otherwise} \end{cases}$

Proof: Jensen, See lemma 2.6 in Tsybakov [Introduction to Nonparametric estimation]

Rk: For instance, gaussian satisfy $KL(N(0, \sigma^2), N(m, \sigma^2)) = \frac{m^2}{2\sigma^2}$

Application to regression:

Say that $X = [-1, 1]^d$, and $d(f, f') = |f(0) - f'(0)|$
 $Y = \mathbb{R}$

$Y = f^*(X) + \varepsilon$ with $\varepsilon \sim N(0, \sigma^2)$

Assume also that $f \in \mathcal{C}^B(L)$: $|f^*(x) - f^*(y)| \leq L|x - y|^B$ $\left\{ \begin{array}{l} X \perp \varepsilon \\ X \sim \text{Unif}([-1, 1]^d) \end{array} \right.$

$\hookrightarrow \mathcal{P}_B := \{d((X, Y), \cdot)\}$

Apply Le Cam's lemma to $P_0 = P_{f_0}$, $P_1 = P_{f_1}$:

(11)

$$R_n(\mathcal{P}) \geq \frac{|f_0(0) - f_1(0)|}{2} \exp(-n KL(P_0, P_1))$$

When $\varepsilon \sim N(0, \sigma^2)$, we can show that $\left. \begin{array}{l} Y_0 = f_0(X) + \varepsilon \\ Y_1 = f_1(X) + \varepsilon \end{array} \right\}$ satisfy $X \sim \text{Unif}([-1, 1]^d)$

$$KL(P_{f_0}, P_{f_1}) \leq \int_{[-1, 1]^d} KL(P_{Y_0|x}, P_{Y_1|x}) dx$$

$$= \int_{[-1, 1]^d} KL(N(f_0(x), \sigma^2), N(f_1(x), \sigma^2)) dx$$

$$= \frac{\|f_0 - f_1\|_{L^2}^2}{2\sigma^2}$$

Hence, $\boxed{R_n(\mathcal{P}) \geq |f_0(0) - f_1(0)| \exp(-n \|f_0 - f_1\|_{L^2([-1, 1]^d)}^2)}$

$\hookrightarrow$ Need to exhibit $f_0, f_1 \in \mathcal{C}^{\beta}(L)$ such that

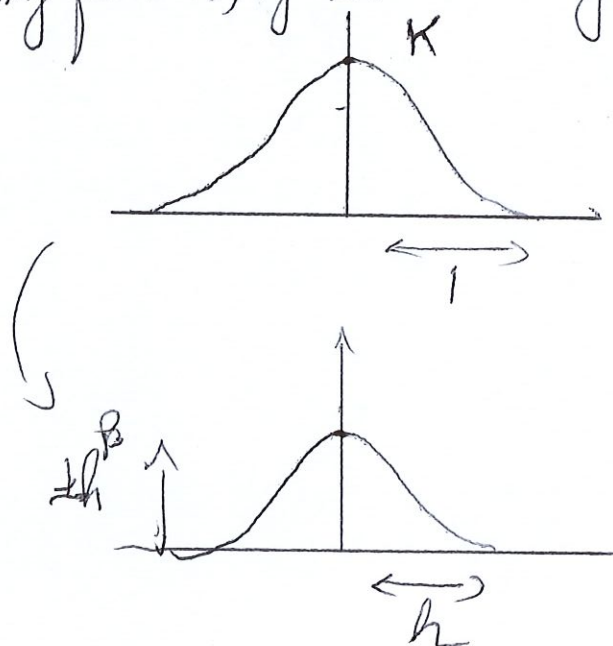
- $|f_0(0) - f_1(0)|$ is large
- $\|f_0 - f_1\|_{L^2([-1, 1]^d)}^2$ is small.

For this, we let $h > 0$ to be chosen later, and

(12)

$$\begin{cases} \cdot f_0 \equiv 0 \\ \cdot f_1(x) = L h^\beta K\left(\frac{x}{h}\right) \end{cases}$$

where $K \in C^p(1)$ is any function, by should be thought of as a localized bump



we have: $\cdot f_0, f_1 \in C^p(L) \forall h > 0$

$$\cdot |f_0(0) - f_1(0)| = L h^\beta K(0)$$

$$\cdot \|f_0 - f_1\|_{L^2}^2 = \int_{B(0,h)} L^2 h^{2\beta} K^2\left(\frac{x}{h}\right) dx = L^2 h^{2\beta+d} \|K\|_{L^2}^2$$

$$\forall h > 0 \Rightarrow \mathcal{R}_m(\mathcal{P}_\beta) \gtrsim L h^\beta \exp\left(-m \frac{L^2 h^{2\beta+d}}{\sigma^2}\right)$$

choosing $\frac{m L^2 h^{2\beta+d}}{\sigma^2} = 1$ (i.e. $h = \left(\frac{\sigma^2}{L^2 m}\right)^{\frac{1}{2\beta+d}}$), we get

$$\mathcal{R}_m(\mathcal{P}_\beta) \gtrsim \left(\frac{\sigma^{-2}}{L^{\frac{d}{p}} m}\right)^{\frac{\beta}{2\beta+d}}$$