

Learning theory from first principles

Chapter 2 : Introduction Decision theory

Data : $(x_i, y_i) \quad i = 1, \dots, n$ training data
 $\downarrow \quad \downarrow$
 $\in X \quad \in Y$

response
label
output

$$\underline{Y \subset \mathbb{R}}$$

Estimate
Compare $f: X \rightarrow Y$

Probability distribution over $(x, y) \in X \times Y$ p
test distribution $\Rightarrow$ training data iid

Machine learning algorithm.

$\mathcal{A}$: data

$(x_1, y_1), \dots, (x_n, y_n) \rightarrow$ function $x \rightarrow y$



Practical evaluation: no access to $p(x, y)$
testing data.

Decision theory

loss function : $\ell : \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}$

$\ell(y, z) =$ loss of predicting z instead of y .

Classification: ($\mathcal{Y}$ finite) $\ell(y, z) = \mathbb{1}_{y \neq z}$

$$\mathcal{Y} = \{1, \dots, k\}$$

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

0-1 loss

Regression: ($\mathcal{Y} = \mathbb{R}$) $\ell(y, z) = (y - z)^2$ square
($|y - z|$ —)

Performance: "surrogates"

Expected risk / test (ing) error / generalization performance

$$\mathcal{R}(f) = \mathbb{E}[\ell(y, f(x))]$$

$f: x \mapsto y$

Expected risk / test (ing) error / generalization performance

$$R(f) = \mathbb{E}[\ell(y, f(x))]$$

f: x → y

Empirical risk / training error

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

$(x_i, y_i) \sim \text{i.i.d. } p(x, y)$

property: if f is deterministic $\mathbb{E}[\hat{R}(f)] = R(f)$

$$\hat{R}(f) - R(f) \approx O\left(\frac{1}{\sqrt{n}}\right)$$

Uniform deviations $\sup_{f \in \mathcal{F}} |\hat{R}(f) - R(f)| = O\left(\frac{1}{\sqrt{n}}\right)$

Examples : Classification $\ell(y, z) = \mathbb{1}_{y \neq z}$

$$R(\ell) = \mathbb{E}[\mathbb{1}_{f(x) \neq y}] = \mathbb{P}(f(x) \neq y)$$

error rate

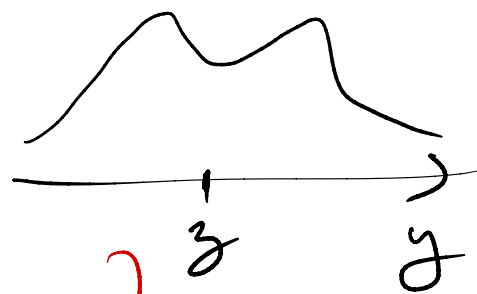
$R(\ell)$ = empirical error rate
proportion of mistakes.

$$f^*(x') = \arg \min_z \mathbb{P}(y \neq z | x = x')$$
$$= \arg \max_z \mathbb{P}(y = z | x = x')$$

Regression : $\ell(y, z) = (y - z)^2$

$$R(\ell) = \mathbb{E}[(y - f(x))^2]$$
$$\bar{R}(\ell) = \frac{1}{n} \sum_{i=1}^n (y_i - f(x_i))^2$$

$$f^*(x') \in \arg \min_{z \in \mathcal{Y}} \mathbb{E}[\ell(y, z) | x = x']$$
$$= \mathbb{E}[y | x = x']$$



Optimal predictor : $f^* = \arg \min_{f \text{ measurable}} R(f)$

Excess risk
 $R(f) - R^*$

= Bayes predictor

$R(f^*) = R^* = \text{Bayes risk}$

$$R(f) = \int \ell(y, f(x)) d\rho(x, y) = \mathbb{E}(\ell(y, f(x)))$$
$$= \mathbb{E}_{x'} \left[\underbrace{\mathbb{E}(\ell(y, f(x')) \mid x = x')}_{\text{depends on } x'} \right]$$

$$f^*(x') \in \arg \min_{z \in \mathcal{Y}} \mathbb{E}(\ell(y, z) \mid x = x').$$

"Optimal ML": given x' , "access" $p(y \mid x = x')$
 $\min_z \mathbb{E}(\ell(y, z) \mid x = x')$

- Two main classes of ML algorithms
- local averaging (h-du)
 - ERM
 - No free lunch
 - Hoeffding.
-

$C^4 43 \rightarrow 11^{400}$.

local averaging

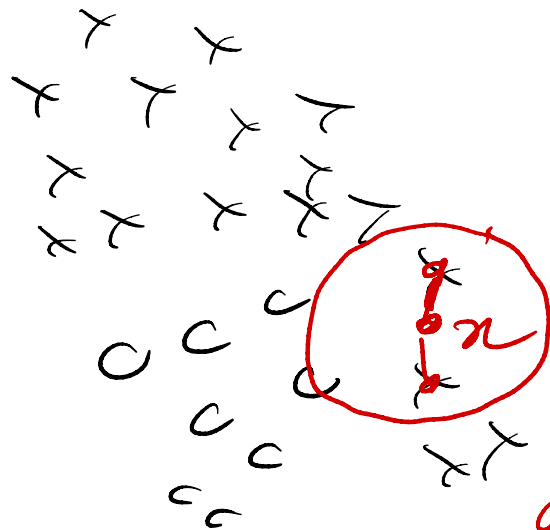
$$f^*(x) = \arg \min_{z \in Y} \int \ell(y, z) d\mu(y|x)$$

↑ replace some estimate

$$d\hat{\mu}(y|x)$$

Curse of dimensionality

$$R(\hat{f}) - R^* \leq O\left(\frac{1}{n^{1/d}}\right)$$



$$X = \mathbb{R}^d$$

• h -nearest neighbors

$$y = \{x_i\}$$

• kernel regression
Watson

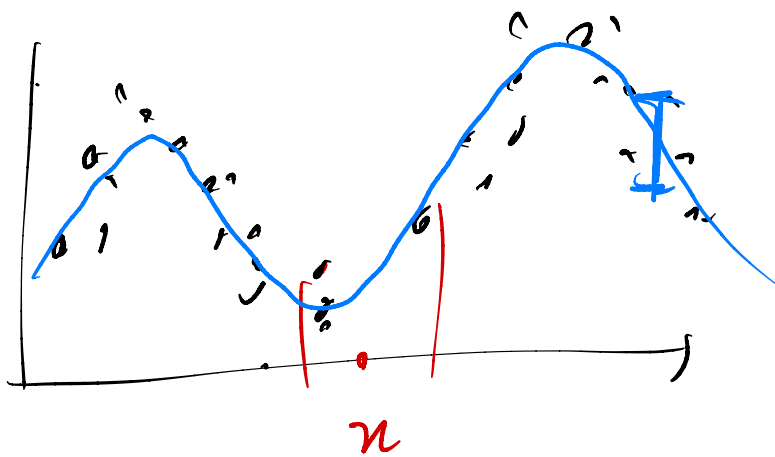
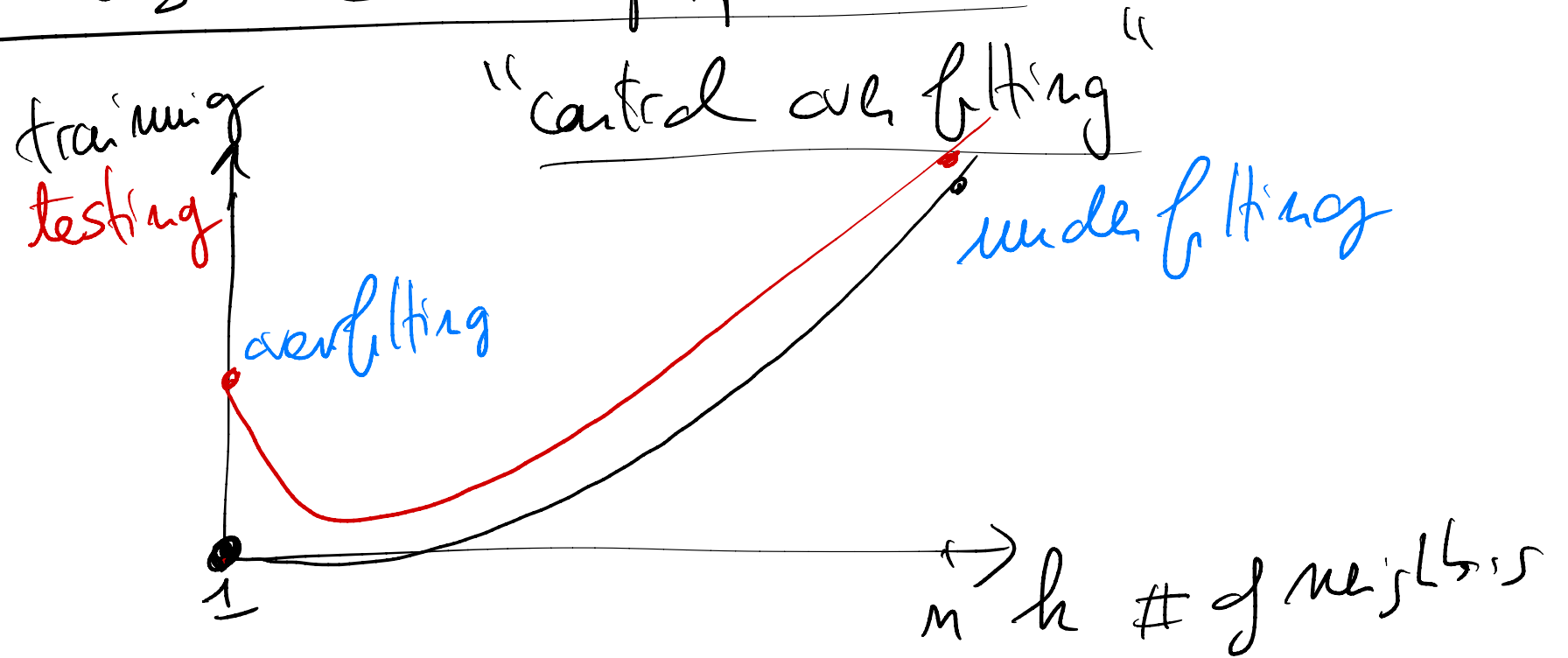
$$d\mu(y|x) = \sum_{i=1}^n w_i(x) \delta_{y=y_i} \text{ Chap 16}$$

$$X = \mathbb{R}$$

$$Y = \mathbb{R}$$



All methods have an hyperparameter



Empirical Risk minimisation ERM

Estimator: model $f_{\theta}: X \rightarrow Y$

θ parameter $\in \Theta$

$$\hat{\theta} = \arg \min_{\theta \in \Theta} \hat{\mathcal{L}}(f_{\theta}) = \frac{1}{n} \sum_{i=1}^n \ell(y_i, f_{\theta}(x_i))$$

$$\hat{f} = f_{\hat{\theta}}$$

$Y = \{-1, 1\} \Rightarrow$ convex surrogates
"logistic loss"

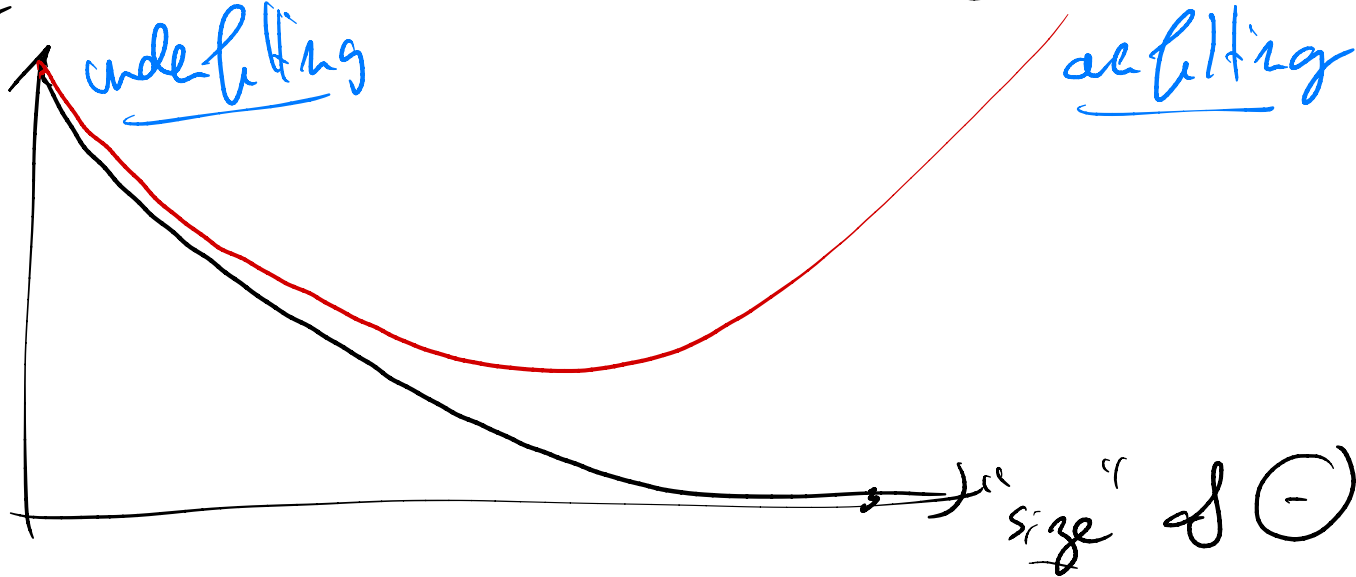
• no practical issues

non-convexity.

training
testing

underfitting

overfitting



Risk decomposition : model $\theta \in \Theta$
 $\hat{\theta} = \arg \min_{\theta \in \Theta} R(\theta)$

expected risk
 $\uparrow$

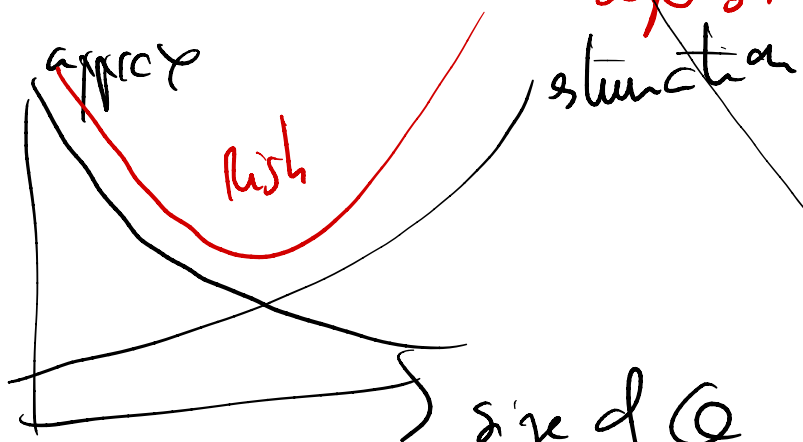
θ' minimize JR

$$\underbrace{R(\hat{\theta}) - R^*}_{\geq 0} = \underbrace{R(\hat{\theta}) - \inf_{\theta' \in \Theta} R(\theta')}_{\text{estimator error}} + \underbrace{\inf_{\theta' \in \Theta} R(\theta') - R^*}_{\text{approximation error}}$$

estimator error
 ≥ 0
 depends on n

approximation error
 ≥ 0

goes down with size of Θ
 deterministic



size of Θ

$$\begin{aligned}
 &= R(\hat{\theta}) - R(\theta') + R(\theta') - R^* \leq 0 \quad \text{opt. error} \\
 &= \underbrace{R(\hat{\theta}) - \hat{R}(\hat{\theta})}_{\text{size of } \Theta} + \underbrace{\hat{R}(\hat{\theta}) - \hat{R}(\theta')}_{\sim \frac{1}{\sqrt{n}}} + \underbrace{\hat{R}(\theta') - R(\theta')}_{\sim \frac{1}{\sqrt{n}}} + \underbrace{R(\theta') - R^*}_{\text{approx. error}}
 \end{aligned}$$

$\leq \sup_{\theta \in \Theta} |R(\theta) - \hat{R}(\theta)|$ uniform deviation

Study of $\hat{Q}(\hat{p}_Q) - \hat{R}(\hat{p}_Q) \leq \hat{Q}(\hat{p}_Q) - \inf_{Q \in \Theta} \hat{Q}(\hat{p}_Q)$

approximate ERM

minimize of $R(\hat{p}_Q)$
 $Q \in \Theta$

No name

$= 0$ if $\hat{Q} = \text{ERM}$
 ≥ 0 otherwise

optimization error

Goals of statistical learning theory

Notion of consistency

Training data $(x_i, y_i), i=1, \dots, n$ IID from distribution $p(x, y)$
 $=: D_n(p)$

ML Algo: $\mathcal{H} : D_n(p) \rightarrow \text{function } x \mapsto y$

$p(x)$ $p(y|x)$
1 1

Good $\mathcal{R}(\mathcal{H}(D_n(p))) - \mathcal{R}_p^*$ small

training error

$\mathcal{R}_p^*$

Bayes risk for p

two types of bounds $\rightarrow$ in high probability : PAC
margin-asymptotic vs. asymptotic probably approx correct

$$(*) \quad \mathbb{E} \left(\mathcal{L}(A(D_n(r))) - \mathcal{L}_p^*$$

Algorithm A is consistent for problem P

$$\text{if } (*) \xrightarrow{n \rightarrow \infty} 0$$

Algorithm A is universally consistent, if $\forall P$ it is consistent

Consistency over a class $\mathcal{P}$ of problems

$$(*) \quad \sup_{P \in \mathcal{P}} \mathbb{E} \left(\mathcal{L}(A(D_n(r))) - \mathcal{L}_P^* \right) \leq \frac{C}{n^\alpha}$$

No free lunch: if $\mathcal{P} = \text{all problems}$

sup(*) arbitrarily slow

- Least-squares regression
- Uniform deviation $\Rightarrow$ all linear models
- Optimization $\begin{cases} \text{GD} \\ \text{SGD} \end{cases}$
- kernel methods
- neural networks

13450