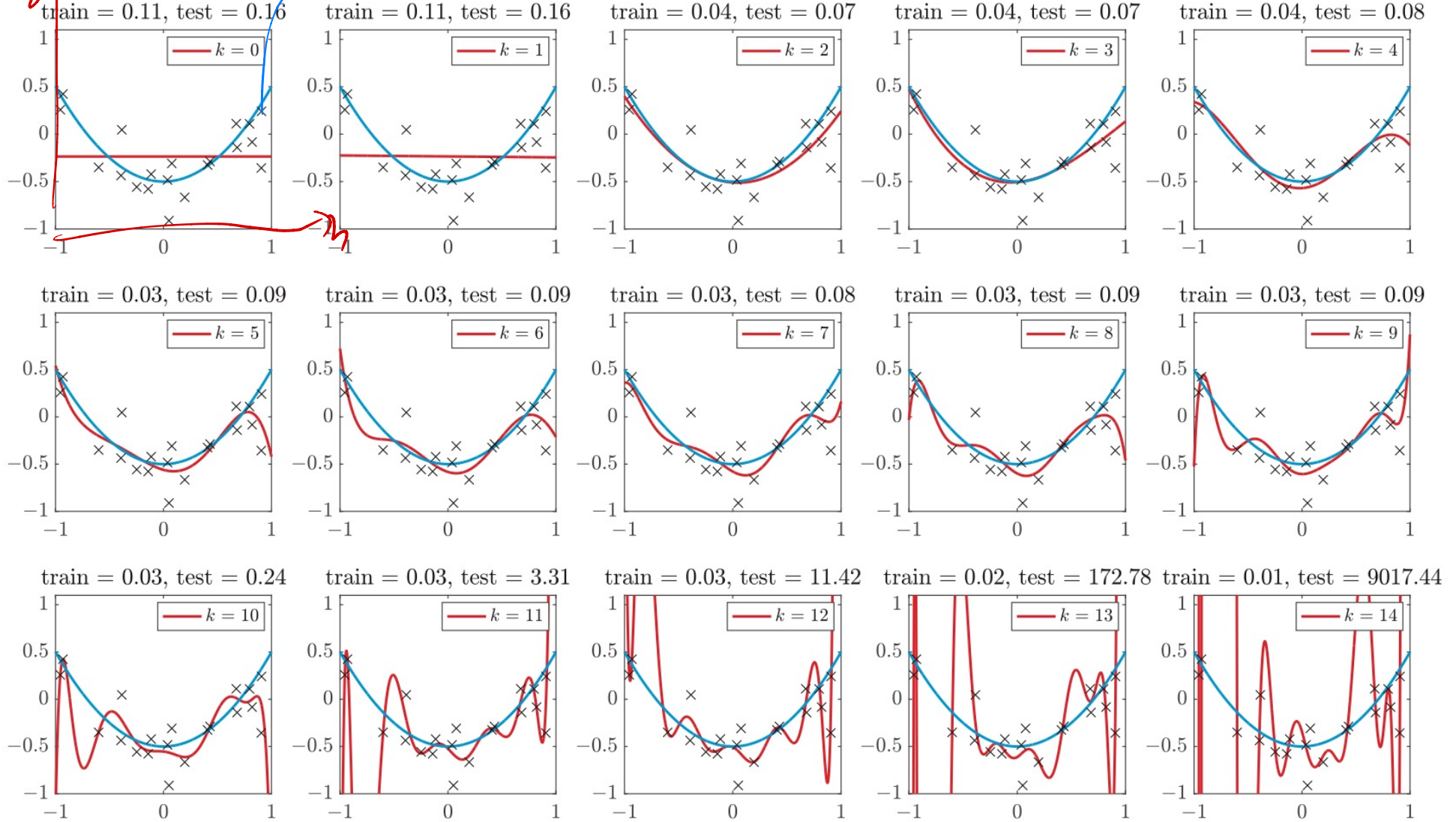


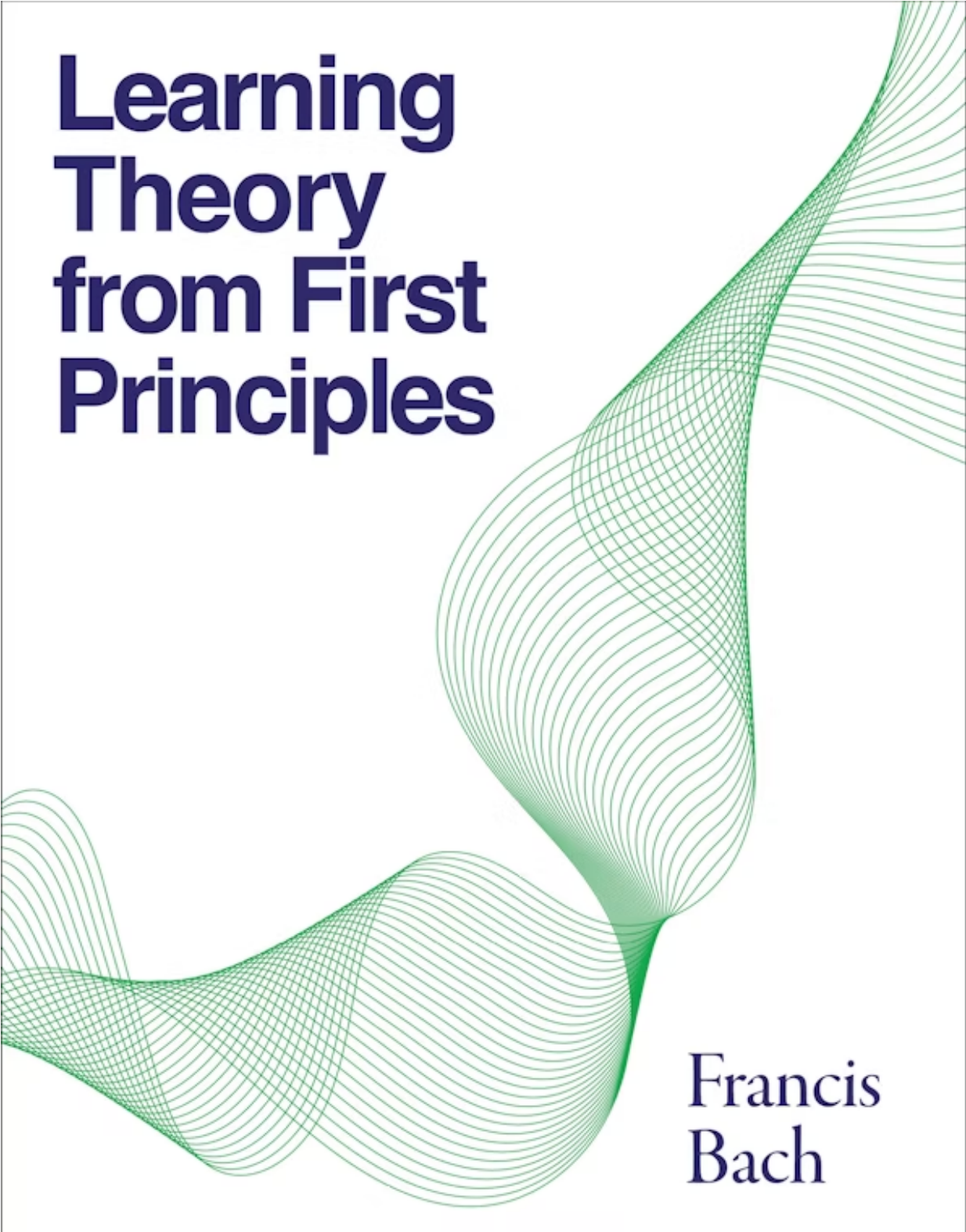
8

$$f^*(x) = \mathbb{E}(y|x)$$

$f_{\theta}(x)$ = polynomial of order n



Learning Theory from First Principles

An abstract graphic consisting of numerous thin, green, wavy lines that overlap and intersect to form a series of flowing, organic shapes. These shapes resemble stylized waves or smoke, moving from the bottom left towards the top right of the page. The lines are closely spaced in some areas, creating a denser, more textured appearance, while in other areas they are more spread out.

Francis
Bach

Supervised learning

Training Data (x_i, y_i) , $i = 1, \dots, n$ — large

$\uparrow$
inputs
covariates
features

outputs
labels
responses

$$x \in \mathcal{X}$$

$$y \in \mathcal{Y}$$

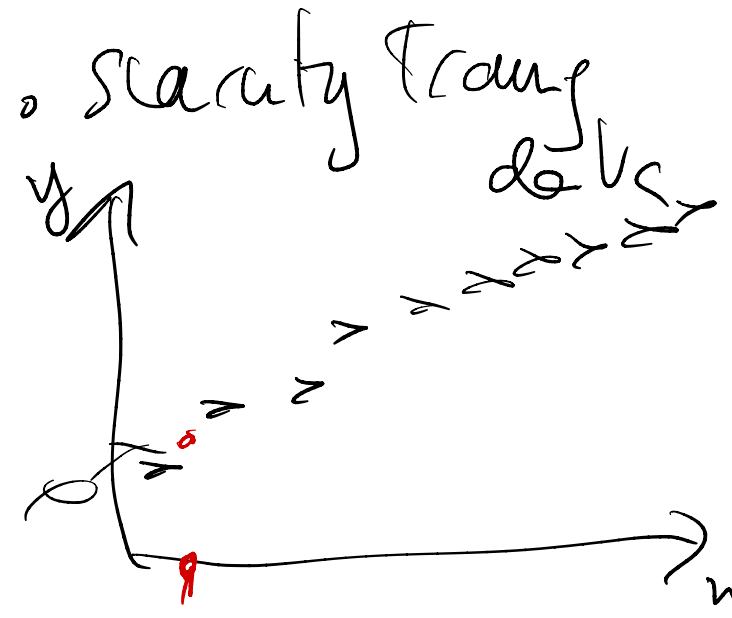
$\hat{\mathbb{R}}^d$

$\phi: \mathcal{X} \rightarrow \mathbb{R}^d$ — large

feature vector

Test performance: $x \in \mathcal{X}$

$$y = f(x) + \text{"noise"}$$



Probabilistic formulation

(x_i, y_i) IID $i=1 \rightarrow n$
independent and identically distributed
from distribution $p(x, y) = \text{test distribution}$

Research: dependence through Markov chain,
domain adaptation

Practical performance evaluation

test data $(x_i^{\text{test}}, y_i^{\text{test}})_{i \in \{1, \dots, n\}}$ — $\sqrt{\frac{\log(\# \text{model})}{n}}$

$$\left(\mathbb{E} \ell(y, f(x)) - \frac{1}{n} \sum_{i=1}^n \ell(y_i^{\text{test}}, f(x_i^{\text{test}})) \right) = O\left(\frac{1}{\sqrt{n}}\right)$$

|-----|
training validation

|-----|
test

Decision theory: Learning when best distribution is known

distribution $p(x, y)$ on $X \times Y$

$$E(f(x)) = \int_X f(x) \underline{d p(x)}$$

$$E(g(x, y)) = \int_{X \times Y} g(x, y) \underline{d p(x, y)}$$

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

$\ell(y, z)$ = loss of predicting z while the observed output is y

• Binary classification: 0-1 loss $\ell(y, z) = \mathbb{1}_{y \neq z}$

$$Y = \{0, 1\}, \quad \boxed{\{-1, 1\}}, \quad \{\pm 2\}$$

• Multiclass: $Y = \{1, \dots, k\}$

• Regression: $Y = \mathbb{R}$ $\ell(y, z) = (y - z)^2$ or $\frac{1}{2} (y - z)^2$

$$\log(1 - e^{-y f(x)})$$

Expected risk
generalization prob
test error
testing error

$$f: \mathcal{X} \rightarrow \mathcal{Y}$$

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

testing data

"functional"

two sources of randomness $\leftarrow$ training data
testing data

Empirical risk
training error

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

training data.

Examples: classification $f: \mathcal{X} \rightarrow \mathcal{Y}$ finite with 0-1 loss
 $R(f) = \mathbb{E}(\mathbb{1}_{y \neq f(x)}) = \mathbb{P}(y \neq f(x))$ error rate

Regression: $R(f) = \mathbb{E}(y - f(x))^2$ "least-squares" $\sim$ accuracy

Bayes Risk and Bayes predictor

$$R^* = R(f^*)$$

↳ optimal $f \sim f^*$

$$R(f) = \mathbb{E}(e(y, f(z))) = \mathbb{E} \left[\underbrace{\mathbb{E}(e(y, f(z)) | n)}_{\text{function of } z} \right]$$

To simplify: $X = \text{finite}$.

$$= \sum_{n' \in X} \underbrace{\mathbb{E}(e(y, f(z')) | n=n')}_{\text{minimize separately}} \cdot p(n=n')$$

$$f: X \rightarrow Y$$

$$f(n') \quad \forall n' \in X$$

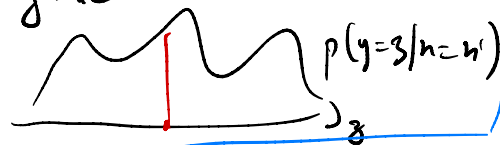
$$f^*(n') \in \arg \min_{z \in Y} \mathbb{E}(e(y, z) | n=n')$$

classifying finite

$$f^*(n') = \arg \min_z P(y \neq z | n=n')$$

$$= \arg \max_{z \in Y} P(y = z | n=n')$$

Regression: $f^*(n') \in \arg \min_{z \in \mathbb{R}} \mathbb{E}(y - z)^2 | n=n')$
 $y = \mathbb{E}(y | n=n')$

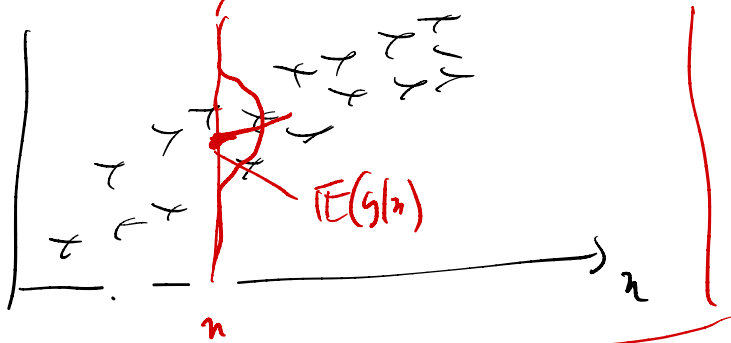


law of total variance
 $\mathbb{E}(y - a)^2 = (\mathbb{E}y - a)^2 + \sigma^2$
 $\sigma^2 = \text{var}(y)$

Lemma: $\arg \min_a \mathbb{E}(y - a)^2 = \mathbb{E}y$

$$\mathbb{E}(y - a)^2 = \mathbb{E}(y^2 - 2ay + a^2) = (\mathbb{E}y)^2 + \sigma^2 - 2a\mathbb{E}y + a^2$$

$$\frac{\partial}{\partial a} = -2\mathbb{E}y + 2a = 0$$



Bayes model: $f^*(x') = \arg \min_{g \in \mathcal{G}} E[\ell(g, y) | x = x']$

Bayes risk = $R(f^*) = E \min_{g \in \mathcal{G}} E[\ell(g, y) | x = x'] = R^*$

Calibration: $R(f) - R^*$ excess risk

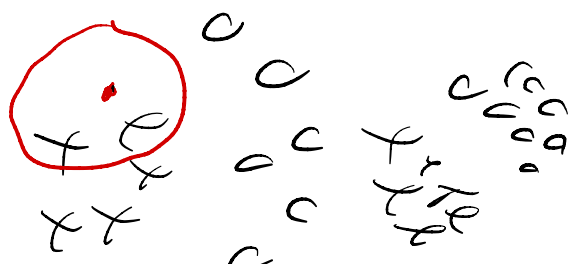
$(c^{4.5} \rightarrow c^4.5)$

Two tests:

- x local averaging
- x Empirical risk min.

local averaging

$$f^*(x) = \operatorname{argmin} \mathbb{E}[l(y, z) | n = \infty] = \int l(y, z) \underbrace{d\eta(y | n = \infty)}_{\downarrow}$$



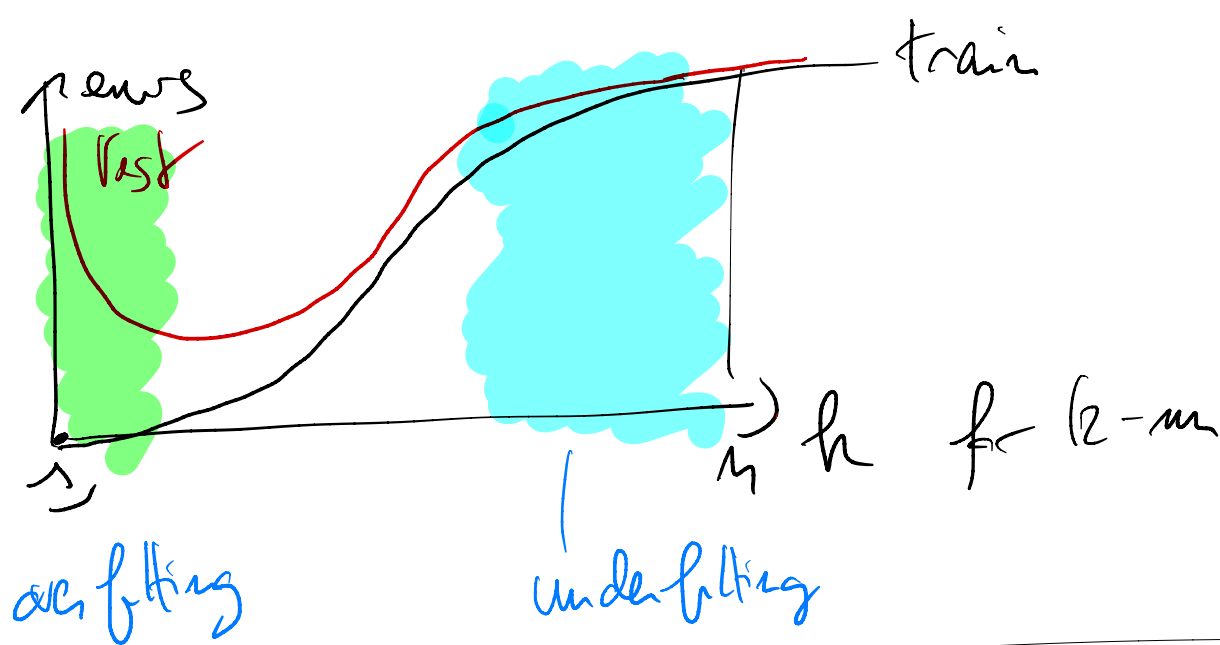
$\hat{d}_p(y | n = \infty)$
training data dependent
kernel

$$\hat{d}_p(y | n) = \sum_{i=1}^n w_i(n) \text{Dirac}(y | y_i)$$

$\rightarrow$ k -nearest neighbors
Nadaraya-Watson

Ex: square loss: $\hat{f}(x) = \sum w_i(n) y_i$

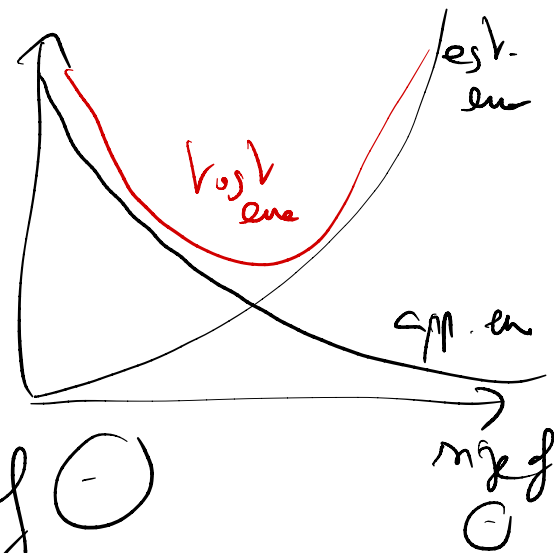
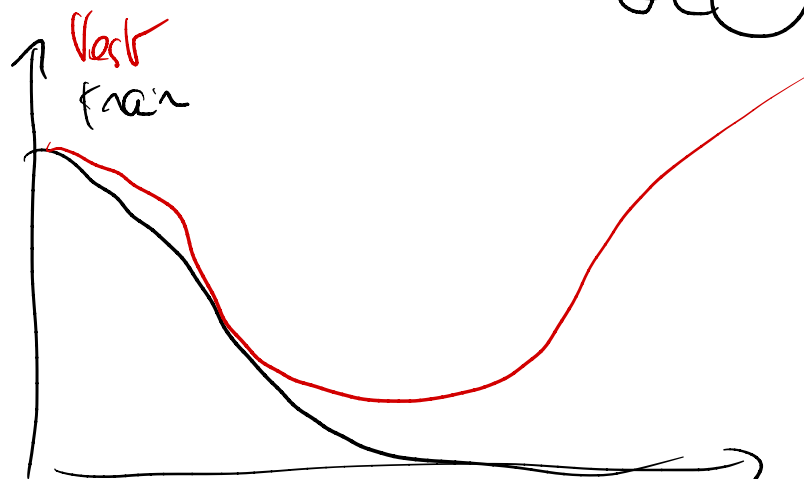
0-1 loss: majority vote weighted by w_i



Empirical Risk Minimization: ERM

Model $f_{\theta}: \mathcal{X} \rightarrow \mathcal{Y}$ of functions parameterized by $\theta \in \Theta$

$$\hat{f} = f_{\hat{\theta}} \text{ s.t. } \hat{\theta} \in \arg \min_{\theta \in \Theta} \hat{\mathcal{R}}(f_{\theta}) = \frac{1}{n} \sum_{i=1}^n \ell(y_i, f_{\theta}(x_i))$$



Decomposition between estimation error & approximation error
 Model $\{f_\theta, \theta \in \Theta\}$.

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

$\geq$
random
 decreasing in
 increasing of Θ

deterministic
 independent of
 decreasing Θ

goal of supervised learning

Data $D_n(p) = (x_i, y_i)_{i=1, \dots, n}$
from dist. p

$\rightarrow$
ML algo
 $\mathcal{A}$

prediction function
 $\hat{f}: x \mapsto y$

$\mathcal{A}$: devises $\rightarrow$ functions

Given $R_p(\hat{f}) - R_p^* = R_p(\mathcal{A}(D_n(p))) - R_p^*$

2 ways of removing randomness $\Rightarrow$ stochastic expectation

$\mathbb{E}_{\text{training data}} R_p(\mathcal{A}(D_n(p))) - R^* \leq \mathcal{R}(p)$

$\hookrightarrow \mathbb{E}_{\text{testing data}}$

$\Rightarrow$ high prob $\forall \delta \in (0, 1)$, with prob $\geq 1 - \delta$
 $R_p(\mathcal{A}(D_n(p))) - R^* \leq \mathcal{R}(\delta)$

PAC

Markov inequality: $X \geq 0$

$$\mathbb{P}(X \geq \varepsilon) \leq \frac{\mathbb{E}X}{\varepsilon} \quad \Leftrightarrow \quad \forall \delta \text{ with prob } 1-\delta, \quad X < \frac{\mathbb{E}X}{\delta}$$

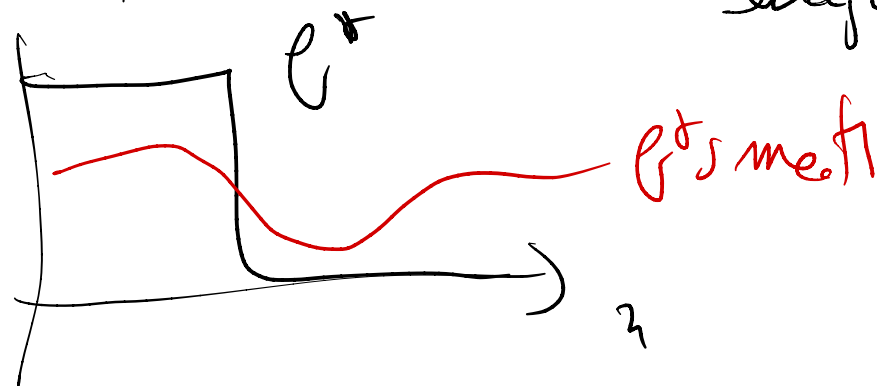
$= \delta$

Sup $\mathbb{P}_{\text{training data}} \mathbb{P}_p(\mathcal{A}(\mathcal{D}_n(p))) - \mathbb{P}_p^*$

$p \in \mathcal{P}$

No Free lunch

$\mathcal{P} = \{g(x, y), \quad \ell^*(x) \text{ smooth, } \text{smooth}\}$



- least-squares
- FEM
- itmp/c6
- local or
kernel methods
- sparse ps
- NN