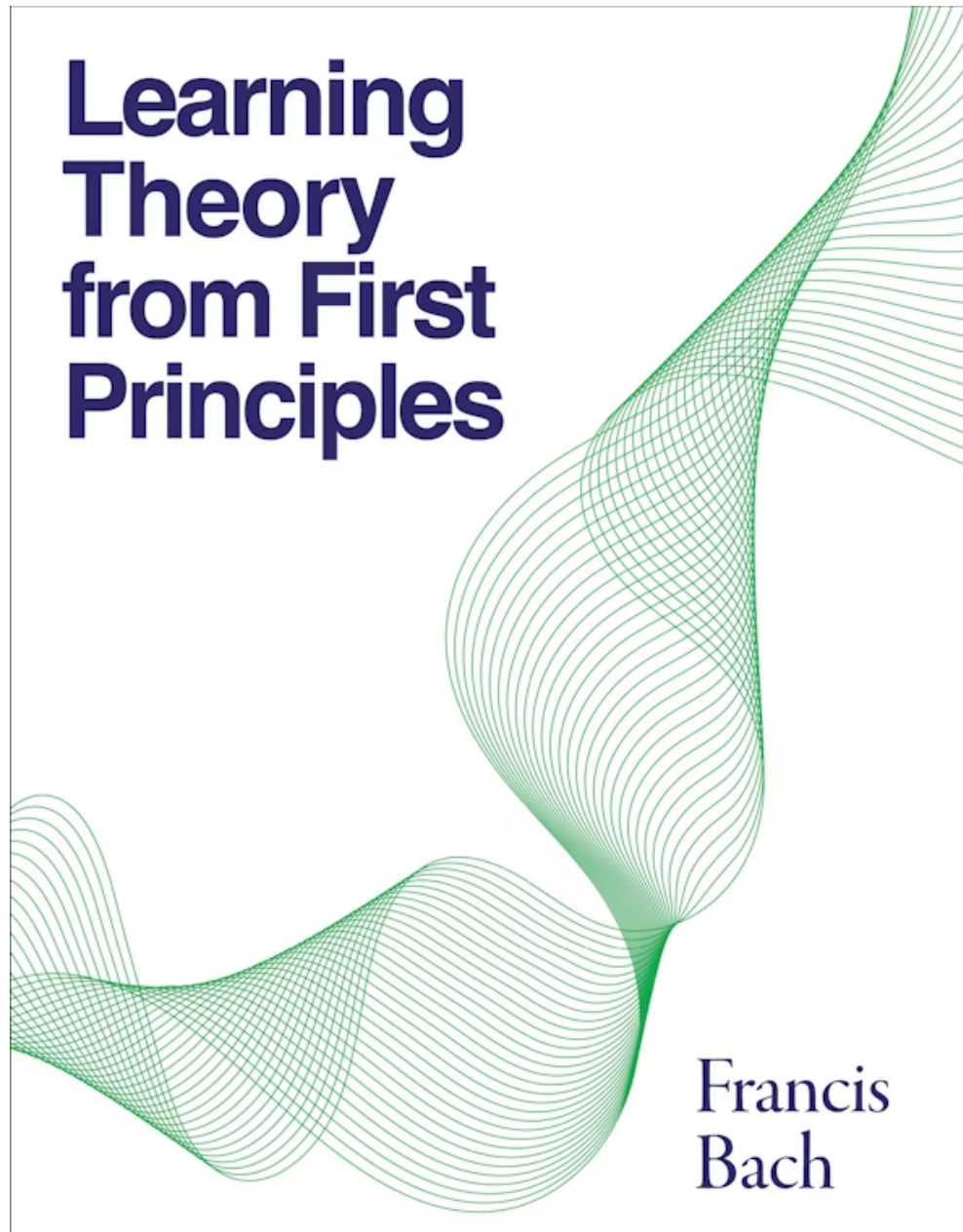


Adaptivity in Machine Learning



- Available online
- Feel free to say!

3 lectures

- ① SGD
- ② kernels
- ③ Neural networks

Simple notes online

Supervised learning

data $(x_i, y_i) \in \mathcal{X} \times \mathcal{Y}$ $i = 1, \dots, n$ IID

Goal: to estimate $f: \mathcal{X} \rightarrow \mathcal{Y}$

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

$\ell(y, z) \rightarrow$ regression $\mathcal{Y} = \mathbb{R}$
 $\ell(y, z) = (y - z)^2$

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

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

Expected risk: $R(\ell) = \mathbb{E}[\ell(y, f(x))]$ $\approx$ test error

function $\mathcal{X} \rightarrow \mathcal{Y}$

art test distribution

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

optimal function f^* $\approx$ Bayes predictor
target function

$$\approx \arg \min R(f)$$

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

Bayes risk

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

$$= \int \left(\int \mathbb{E}[l(y, f(x)) | x] p(y|x) dy \right) p(x) dx$$

$\mathbb{E}(l(y, f(x)) | x)$

$$\Rightarrow f^*(x) = \arg \min_{f \in \mathcal{F}} \mathbb{E}(l(y, f(x)) | x)$$

square loss $\mathbb{E}[y|x]$

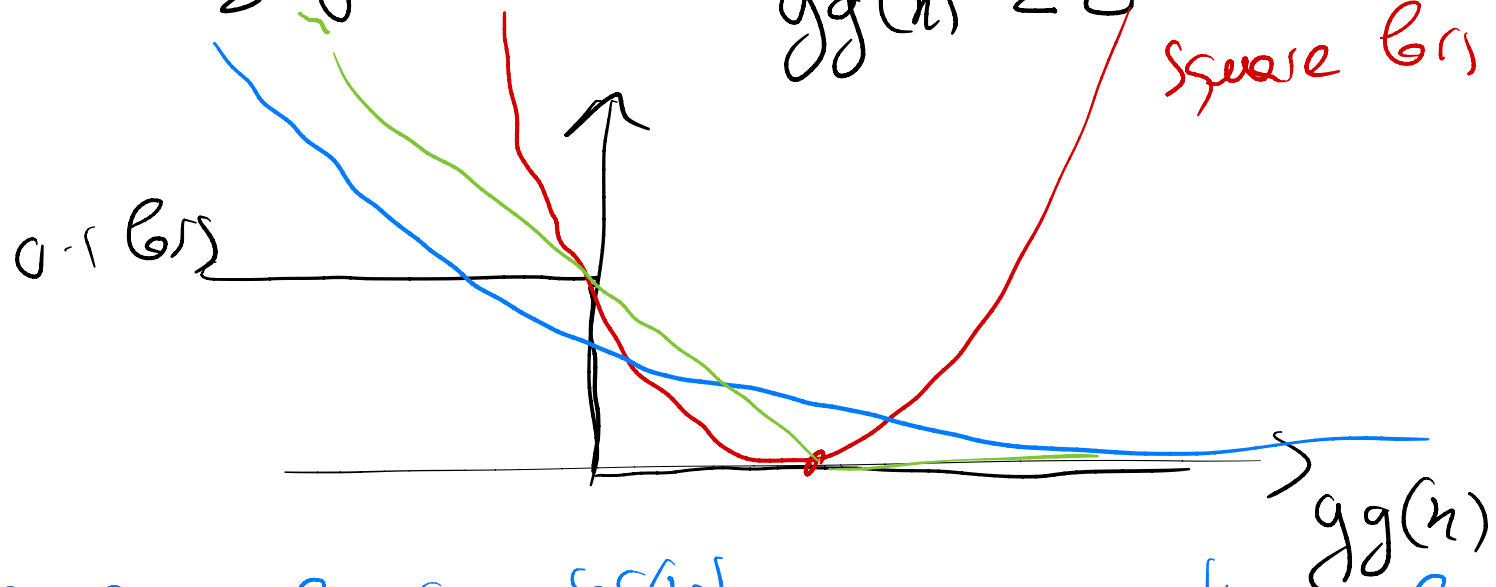


$$0-1 \text{ loss} = f^*(x) = \text{sign}(\mathbb{E}(y|x))$$

$y|x$

Convex surrogates for classification
 instead of estimating $f: \mathcal{X} \rightarrow \{-1, 1\}$
 we estimate $g: \mathcal{X} \rightarrow \mathbb{R}$ and $f = \text{sign}(g)$

0-1 loss of $f(x) = 1$ $gg(x) < 0$



square loss: $(yg(x) - 1)^2$
 $(g(x) - y)^2$

logistic loss: $\log(1 + e^{-yg(x)}) \approx$ cross entropy loss

hinge loss: $(1 - yg(x))_+$

$\Rightarrow$ ASSUME: $\ell(g, y)$ is G -lipschitz
 continuous in y

Analysis: $\mathcal{F} = \{f: x \rightarrow \mathbb{R}\} = \{f_\theta: x \rightarrow \mathbb{R}, \theta \in \Theta\}$

Goal: find $\hat{f} \in \mathcal{F}$ st $R(\hat{f}) - R_*$ small.

decomposition: $R(\hat{f}) - R_* = \underbrace{R(\hat{f}) - \inf_{f \in \mathcal{F}} R(f)}_{\text{estimation error}} + \underbrace{\inf_{f \in \mathcal{F}} R(f) - R_*}_{\text{approximation error}}$

Empirical risk minimizer
 $\hat{f} \approx \arg \min_{f \in \mathcal{F}} \hat{R}(f)$

Estimation error = $\underbrace{R(\hat{f}) - \hat{R}(\hat{f})}_{\geq 0} + \underbrace{\hat{R}(\hat{f}) - \hat{R}(f_F^*)}_{\leq \hat{R}(\hat{f}) - \inf_{f \in \mathcal{F}} \hat{R}(f)} + \underbrace{\hat{R}(f_F^*) - R(f_F^*)}_{\text{approximation error}}$

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

$\mathbb{E} \hat{R}(f) = R(f)$ if f fixed

$\rightarrow \leq \sup_{f \in \mathcal{F}} R(f) - \hat{R}(f)$ uniform deviation

gradient descent : $F: \mathbb{R}^d \rightarrow \mathbb{R}$ convex, B -Lipschitz continuous

ex: (1) $F(\theta) = R(b_\theta)$ with $b_\theta(x) = \theta^T \phi(x)$ where $\phi: \mathcal{X} \rightarrow \mathbb{R}^d$
feature map

convex loss G -Lipschitz continuous

and $\|\phi(x)\| \leq R$ for all $x \Rightarrow B = GR$
Convex

(2) $F(\theta) = R(b_\theta)$

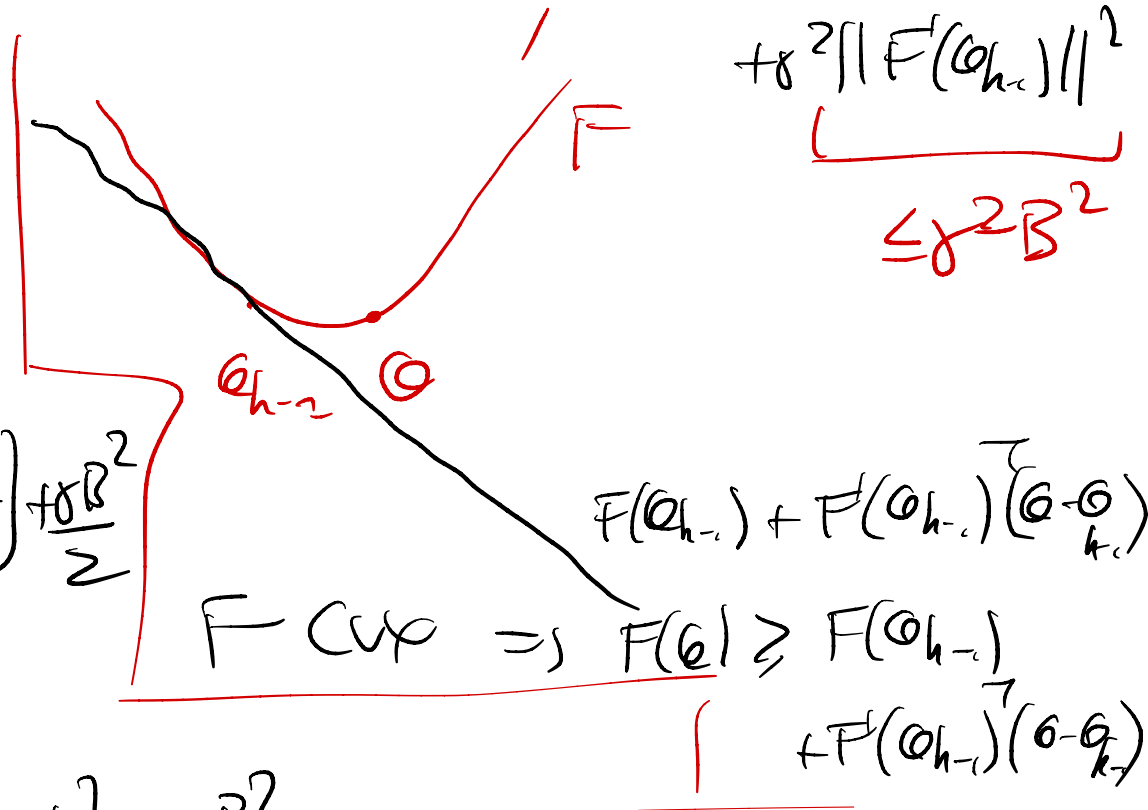
gradient descent = iterative alg. $\theta_n = \theta_{n-1} - \gamma F'(\theta_{n-1})$



Analysis: $Q_k = \Pi_D(Q_{k-1} - \gamma F'(Q_{k-1}))$ project a over ball of center c and radius D

For any $Q \in \mathbb{R}^d$, $\|Q\| \leq D$

$$\begin{aligned} \|Q_k - Q\|^2 &\leq \|Q_{k-1} - \gamma F'(Q_{k-1}) - Q\|^2 = \|Q_{k-1} - Q\|^2 - 2\gamma F'(Q_{k-1})^\top (Q - Q_{k-1}) \\ &\quad + \gamma^2 \|F'(Q_{k-1})\|^2 \\ &\leq \|Q_{k-1} - Q\|^2 + \gamma^2 B^2 - 2\gamma [F(Q_{k-1}) - F(Q)] \end{aligned}$$



$$\Rightarrow F(Q_{k-1}) - F(Q)$$

$$\leq \frac{1}{2\gamma} [\|Q_{k-1} - Q\|^2 - \|Q_k - Q\|^2] + \frac{\gamma B^2}{2}$$

averaging from $k=1$ to t

$$\frac{1}{t} \sum_{k=1}^t F(Q_{k-1}) - F(Q) \leq \frac{1}{2\gamma t} \|Q_0 - Q\|^2 + \frac{\gamma B^2}{2}$$

(Jensen's

$$F\left(\frac{1}{t} \sum_{k=1}^t Q_{k-1}\right) \leq F(Q)$$

$$\leq F(Q) + \frac{1}{2\gamma t} \|Q_0 - Q\|^2 + \frac{\gamma B^2}{2}$$

$$\leq \boxed{\frac{BD}{\sqrt{t}}} + \frac{1}{\sqrt{t}}$$

Applied to ML . $F(\theta) = \frac{1}{n} \sum_{i=1}^n \ell(y_i, \theta^T \phi(x_i))$

$$F'(\theta) = \frac{1}{n} \sum_{i=1}^n \ell'(y_i, \theta^T \phi(x_i)) \phi(x_i)$$

kernels

$$F(\text{averaged kernel}) - F(\theta_0) \leq \frac{GRD}{\sqrt{t}}$$

NB : uniform bounds $\frac{GRD}{\sqrt{n}}$

$\Rightarrow$ optimisation is "free"
eg $t=n$

directions

$\Rightarrow O(n^2)$ scores
to gradients

SGD = minimize $F(\theta) = \mathbb{E} \ell(y, \theta^T \phi(x))$
access to unbiased gradients

$$\theta_h = \theta_{h-1} - \gamma g_h$$

$$\mathbb{E}(g_h | \theta_{h-1}) = F'(\theta_{h-1})$$

stochastic gradient

$$\Rightarrow g_h = \frac{\partial}{\partial \theta} \ell(y_h, \theta^T \phi(x_h))$$

s.t. $\theta = \theta_{h-1}$

F =
test error

Prop

$$\mathbb{E} F\left(\frac{1}{n} \sum_{h=1}^n \theta_h\right) \leq F(\theta) + \frac{G^2 D}{\sqrt{n}}$$

↳ training data

$$\text{if } h \leq n$$

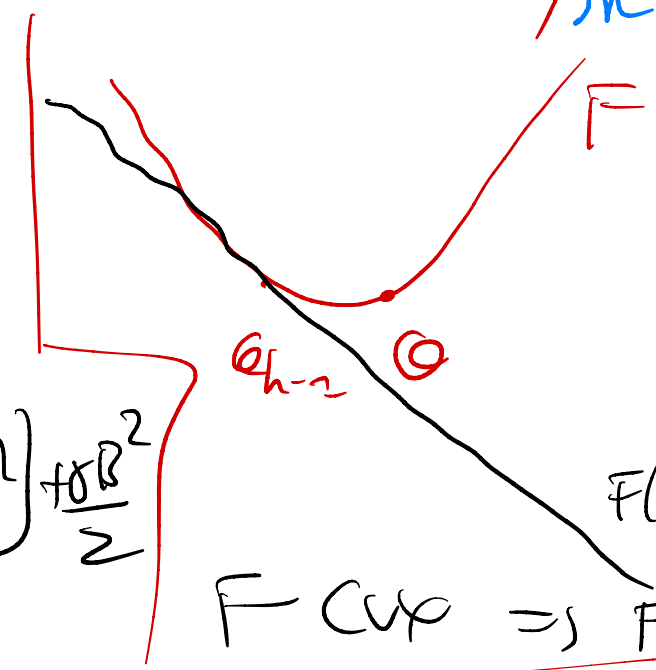
complexity: $O(n)$
access to grad.

Analysis: $Q_k = \Pi_D(Q_{k-1} - \gamma \frac{F'(Q_{k-1})}{g_k})$ project a over ball of center c and radius D

For any $Q \in \mathbb{R}^d$, $\|Q\| \leq D$

$$\|Q_k - Q\|^2 \leq \|Q_{k-1} - \gamma F'(Q_{k-1}) - Q\|^2 = \|Q_{k-1} - Q\|^2 - 2\gamma F'(Q_{k-1})^\top (Q - Q_{k-1}) + \gamma^2 \|F'(Q_{k-1})\|^2$$

$$\leq \|Q_{k-1} - Q\|^2 + \gamma^2 B^2 - 2\gamma [F(Q_{k-1}) - F(Q)]$$



$$= F(Q_{k-1}) - F(Q) \leq \frac{1}{2\gamma} [\|Q_{k-1} - Q\|^2 + \gamma B^2]$$

averaging from $k=1$ to t

$$\frac{1}{t} \sum_{k=1}^t F(Q_{k-1}) - F(Q) \leq \frac{1}{2\gamma t} \|Q_0 - Q\|^2 + \frac{\gamma B^2}{2}$$

(Jensen's)

$$EF\left(\frac{1}{t} \sum_{k=1}^t Q_{k-1}\right) \leq F(Q) + \frac{1}{2\gamma t} \|Q_0 - Q\|^2 + \frac{\gamma B^2}{2} \leq \sqrt{\frac{BD}{\sqrt{t}}} + \frac{1}{\sqrt{t}}$$