

Maths of Machine Learning Summary

CHAPTER ONE

Statistical Learning Theory

joint prob measure P_0

① Classification & Regression: $(X, Y) \in \mathcal{X} \times \mathcal{Y}$

- Goal: learn $h: \mathcal{X} \rightarrow \mathcal{Y}$ to predict Y from X ,
feature response
- Loss function: $\ell: \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}$,
- Risk: choose h to minimise the risk

$$R(h) = \int_{(x,y) \in \mathcal{X} \times \mathcal{Y}} \ell(h(x), y) dP_0(x, y) = \mathbb{E}[\ell(h(X), Y)]$$

	Classification	Regression
Y	categorical	$\mathbb{R}$
$\ell(h(x), y)$	$\mathbb{1}(h(x) \neq y)$	$(h(x) - y)^2$
$R(h)$	$P(h(X) \neq Y)$	$h(x) = \mathbb{E}[Y X=x]$

$\eta(x) = P(Y=1|X=x)$

maximum a-posteriori estimator

② The Bayes classifier: $Y = \{0, 1\}$, characterise $h = \arg\min_h P(h(X) \neq Y)$

- Prop 1: The Bayes classifier $h_*(x) = \arg\min_h \mathbb{E}[\ell(h(x), Y)] = \begin{cases} 1 & \eta(x) > \frac{1}{2} \\ 0 & \eta(x) < \frac{1}{2} \end{cases}$

Pf: $R(h) = \mathbb{E}[\mathbb{1}\{h(X) \neq Y\}] = \mathbb{E}[\mathbb{E}[\mathbb{1}\{h(X) \neq Y\} | X]]$ [Tower property]

(*) $= \mathbb{E}[\mathbb{1}\{h(X)=0, Y=1\} + \mathbb{1}\{h(X)=1, Y=0\} | X]$ [cases on $\neq$]

$= \mathbb{1}\{h(X)=0\} \underbrace{\mathbb{E}[Y=1|X]}_{\eta(x)} + \mathbb{1}\{h(X)=1\} \underbrace{\mathbb{E}[Y=0|X]}_{1-\eta(x)}$ [taking out known]

If $\eta(x) > 1-\eta(x) \Rightarrow h(x)=0$ minimises
 If $\eta(x) < 1-\eta(x) \Rightarrow h(x)=1$ minimises } cases $\Rightarrow h = h_*$

Motivation: in practice, we observe data $\{(x_i, y_i)\}_{i=1}^n$, w/ P_0 probability distribution unknown!
 Instead of estimating P_0 from data, we'd desire $\mathcal{H}$ of piece best $h \in \mathcal{H}$.

e.g. $\mathcal{H} = \{h: h(x) = \text{sgn}(b + \sum_{j=1}^n w_j \phi_j(x))\}$, $w \in \mathbb{R}^d, b \in \mathbb{R}$
Given set of basic functions $\phi_j: \mathcal{X} \rightarrow \mathbb{R}$



③ Empirical Risk Minimisation & Hypothesis classes

- Suppose $(X_i, Y_i) \stackrel{iid}{\sim} P_0$ for $i \in \{1, \dots, n\}$. The empirical risk is

$$\hat{R}(h) = \frac{1}{n} \sum_{i=1}^n \ell(h(X_i), Y_i), \quad \hat{h} = \arg\min_{h \in \mathcal{H}} \hat{R}(h), \quad \text{clearly } \mathbb{E}[\hat{R}(h)] = R(h)$$

- E.g. Linear regression

- $\mathcal{X} = \mathbb{R}^p, \mathcal{Y} = \mathbb{R}, \ell(z, y) = (z - y)^2, \mathcal{H} = \{x \mapsto b + x^T w, w \in \mathbb{R}^p, b \in \mathbb{R}\}$

- $\hat{h}(x) = x^T \hat{w} + \hat{b}$ w/ $(\hat{b}, \hat{w}) = \arg\min_{w, b} \frac{1}{n} \sum_{i=1}^n (Y_i - x_i^T w - b)^2$

Pf: $\mathbb{E}[\hat{R}(h)] = \mathbb{E}[\frac{1}{n} \sum_{i=1}^n \ell(h(X_i), Y_i)] = \frac{1}{n} \sum_{i=1}^n \mathbb{E}[\ell(h(X_i), Y_i)] = \frac{1}{n} \cdot n \cdot R(h) = R(h)$

- E.g. k-nearest neighbours

- $\mathcal{X} = \mathbb{S}^{d-1}$, consider k nearest training pts to $\bar{x}$ and assign the most common label.

- $d_{\bar{x}} = (\| \bar{x} - x_i \|)_{i=1}^n$, $I_{\bar{x}}$ indexes k smallest entries in $d_{\bar{x}}$, $h(\bar{x}) = \text{sgn}(\frac{1}{k} \sum_{i \in I_{\bar{x}}} Y_i)$

n-vector of distances

larger hypothesis class better approximates

vote on class by k-nearest nbs

④ Bias-Variance Trade off: larger $\mathcal{H} \Rightarrow$ best h_* estimator, but more sensitivity to initial data.

- Regression Example

- Dataset $D = \{(x_i, y_i)\}_{i=1}^n$ drawn iid from P_0 , Note: cross term disappears

- $\bar{y}(x) = \mathbb{E}[Y|X=x]$ expected output

$\mathbb{E}[\mathbb{E}[(h_0(x) - \bar{y}(x))(\bar{y}(x) - Y)|X]] = 0$

- $h_0(x)$ is the hypothesis learnt from D

- $\bar{h}(x) := \mathbb{E}_{D \sim P_0^n}[h_0(x)]$ expected classifier by tower property & taking out known

expected test error

$\mathbb{E}[(h_0(x) - Y)^2] = \mathbb{E}[(h_0(x) - \bar{y}(x) + \bar{y}(x) - Y)^2] = \underbrace{\mathbb{E}[(h_0(x) - \bar{y}(x))^2]}_{(*)} + \underbrace{\mathbb{E}[(\bar{y}(x) - Y)^2]}_{=0}$

(*) $= \mathbb{E}[(h_0(x) - \bar{h}(x) + \bar{h}(x) - \bar{y}(x))^2] = \underbrace{\mathbb{E}[(h_0(x) - \bar{h}(x))^2]}_{\text{variance}} + \underbrace{\mathbb{E}[(\bar{h}(x) - \bar{y}(x))^2]}_{\text{bias}} + 2\mathbb{E}[(h_0(x) - \bar{h}(x))(\bar{h}(x) - \bar{y}(x))]$

$\mathbb{E}_0(h_0(x)) = \bar{h}(x)$

- Notes do same w/ k-NN
[skipped as too hard & slow]

$\mathcal{H} \uparrow$ complex $\Rightarrow \uparrow$ variance, $\mathcal{H} \downarrow$ bias

- Cross Validation: Split dataset $D = \{(x_i, y_i)\}_{i=1}^n$ into D_{train} and D_{test} .
Pick the estimator which minimizes the empirical risk.
→ V -fold splits data into V sets, pick $\hat{h}_V = \argmin_h \frac{1}{V} \sum_{i=1}^V E_{i, \text{tr}}$

⑥ Excess Risk: Given an estimator $\hat{h}$, the excess risk is given by

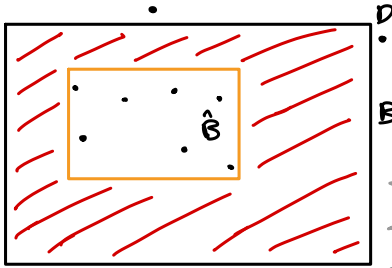
$$\mathcal{E}(\hat{h}) = R(\hat{h}) - R(h_*)$$

Bayes classifier "best you can do"

Note: "No free lunch" theorem $\Rightarrow$ impossible to learn all probability distributions w/out some assumptions.

Idea: will theoretically analyze this for different classes:
- How does the complexity affect $\mathcal{E}$?
- How does # data pts affect $\mathcal{E}$?

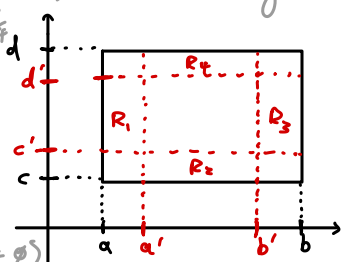
- Example (Learning Rectangles): $X = \mathbb{R}^2$, $Y = \{0, 1\}$, goal: learn $h(x) = \mathbb{1}_B(x) = \begin{cases} 1 & \text{if } x \in B \\ 0 & \text{if } x \notin B \end{cases}$
Given iid data $(x_i, y_i) \in X \times Y$ w/ $y_i = \begin{cases} 1 & \text{if } x_i \in B \\ 0 & \text{if } x_i \notin B \end{cases}$, $\hat{h} = \mathbb{1}_{\hat{B}}$ where $\hat{B}$ is the smallest rectangle containing pts labeled as inside B . Note: $R(h_*) = 0 \therefore h_* = \mathbb{1}_B$ is the Bayes classifier
 $\mathcal{E}(\hat{h}) = R(\hat{h}) = P(\hat{h}(x) \neq \mathbb{1}_B(x)) = P(X \in B \setminus \hat{B}) + P(X \in \hat{B} \setminus B)$
Note: $R(h_*) = 0 \therefore \hat{B} \subset B$



Q: How many samples to ensure $P(R(\hat{h}) \leq \epsilon) \geq 1 - \delta$?

Assume $P(X \in B) > \epsilon$ o/w, $R(\hat{h}) = P(X \in B \setminus \hat{B}) \leq \epsilon$ trivially!

- Define $R_1 = [a, a'] \times [c, d]$, a' smallest s.t. $P(X \in R_1) \geq \epsilon/4$
- Def $R_2 = [a, b] \times [c, c']$, c' " s.t. $P(X \in R_2) \geq \epsilon/4$
- Def $R_3 = [b', b] \times [c, d]$, b' largest s.t. $P(X \in R_3) \geq \epsilon/4$
- Def $R_4 = [a, b] \times [d', d]$, d' largest s.t. $P(X \in R_4) \geq \epsilon/4$



Note that

$$\{ \hat{B} \cap R_i \neq \emptyset \forall i \} \Rightarrow P(X \in B \setminus \hat{B}) \leq \epsilon \text{ so } P(R(\hat{h}) \leq \epsilon) \geq P(\hat{B} \cap R_i \neq \emptyset \forall i) \\ = 1 - P(\exists i \text{ s.t. } \hat{B} \cap R_i = \emptyset) \geq 1 - \sum_{i=1}^4 P(\hat{B} \cap R_i = \emptyset) \quad [\text{union bd}]$$

$$P(\hat{B} \cap R_i = \emptyset) = \prod_{j=1}^n P(x_j \notin R_i) \leq (1 - \epsilon/4)^n \leq e^{-\frac{\epsilon n}{4}} \text{ so can put together}$$

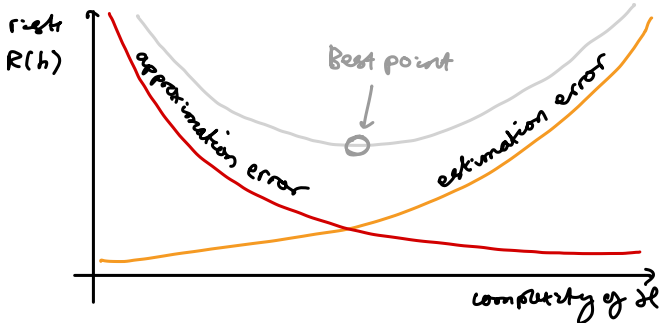
$$P(R(\hat{h}) \leq \epsilon) \geq 1 - 4e^{-\frac{\epsilon n}{4}} \geq 1 - \delta \Leftrightarrow \delta > 4e^{-\frac{\epsilon n}{4}} \Leftrightarrow -\log(\frac{\delta}{4}) \geq -\frac{\epsilon n}{4} \Leftrightarrow n \geq \frac{4}{\epsilon} \log(\frac{4}{\delta})$$

Note: As $h_* = \mathbb{1}_B$, $R(h_*) = 0$, now consider $h_* \notin \mathcal{H}$ so $R(h_*) > 0$

we want to minimize the excess risk

excess risk

- Decomposition of Excess Risk: $R(\hat{h}) - R(h_*) = \underbrace{R(\hat{h}) - \inf_{h \in \mathcal{H}} R(h)}_{\text{Estimation Error [variance]}} + \underbrace{\inf_{h \in \mathcal{H}} R(h) - R(h_*)}_{\text{Approximation error [Bias]}}$



- Approximation error: Need assumptions on h_* & $\mathcal{H}$.

e.g. $z \mapsto \varphi(y, z)$ L -Lipschitz w.r.t. z . Then

$$R(h) - R(h') = \mathbb{E}[\varphi(h(x), y) - \varphi(h'(x), y)] \leq \mathbb{E}[|h(x) - h'(x)|]$$

Assume $h_*(x) = w_*^T \varphi(x)$ for $\varphi(x) = (1, x_1, \dots, x_p, x_1^2, x_1^3, \dots)$

and $\mathcal{H} = \{h_w(x) = w^T \varphi(x); \|w\| \leq k\}$. Then

$$\inf_{h \in \mathcal{H}} R(h) - R(h_*) \leq L \inf_{\|w\| \leq k} \mathbb{E}[|\varphi(x), w - w_*|] \stackrel{\text{zero}}{\leq} L \mathbb{E} \|\varphi(x)\| \max_{\|w\| \leq k} \|w - w_*\|$$

$$\leq L \mathbb{E} \|\varphi(x)\| \max_{\|w\| \leq k} \|w - w_*\|$$

$$= \inf_{\|w\| \leq k} \mathbb{E} \|w - w_*\|$$

- ⑦ The estimation Error: consider for ERM $\hat{h}$.

letting $\bar{h} = \argmin_{h \in \mathcal{H}} R(h)$, note that $\hat{R}(\hat{h})$ is as small as possible $\therefore \hat{h}$ is ERM

$$R(\hat{h}) - R(\bar{h}) = \underbrace{R(\hat{h}) - \hat{R}(\hat{h})}_{\text{estimation error for } \hat{h}} + \hat{R}(\hat{h}) - \hat{R}(\bar{h}) + \hat{R}(\bar{h}) - R(\bar{h}) (*)$$

$$\leq \sup_{h \in \mathcal{H}} R(h) - \hat{R}(\hat{h}) + \hat{R}(\bar{h}) - R(\bar{h}) \quad [\text{relaxing conditions}]$$

Note: CLT is asymptotic result so no quantitative bounds... we need concentration inequalities!

$$\leq 2 \sup_{h \in \mathcal{H}} |R(h) - \hat{R}(h)|$$

$$\text{GOAL: } P(\sup_{h \in \mathcal{H}} |R(h) - \hat{R}(h)| > \epsilon) \leq \delta$$

Precisely control the estimation error of our hypothesis class $\mathcal{H}$

Tools from Probability

$$P\left(\sum_{i=1}^n x_i - E[x_i] \geq \delta\right) \leq \exp\left(-\frac{2\delta^2}{\sum_{i=1}^n (b_i - a_i)^2}\right) \quad [\text{HDP}]$$

① Markov: $P(W \geq t) \leq \frac{E[W]}{t} \Rightarrow P\left(\frac{1}{n} \sum_{i=1}^n (x_i - E[x_i]) \geq t\right) \Rightarrow \text{set } \delta = tn \text{ hence } n^2$

② Chernoff: $P(W \geq t) \leq \inf_{\alpha > 0} e^{-\alpha t} E[\exp(\alpha W)]$

③ Sub-Gaussian: W sub-gaussian w/ param $\sigma \Rightarrow E[e^{\alpha(W - E[W])}] \leq e^{\alpha^2 \sigma^2 / 2} \quad \forall \alpha \in \mathbb{R}$

$\hookrightarrow$ If W sub-gaussian, $P(W - E[W] \geq t) \leq e^{-t^2 / 2\sigma^2} \quad \forall t \geq 0$

$\hookrightarrow$ If $W \in [a, b] \Rightarrow W$ sub-gaussian w/ $\sigma = \frac{b-a}{2}$

$\hookrightarrow W_1, \dots, W_n$ independent. W_i sub-gaussian w/ param σ_i . Then $\forall t \in \mathbb{R}$,

$$\sum_{i=1}^n W_i \text{ sub-g. w/ param } \left(\sum_{i=1}^n \sigma_i^2\right)^{1/2}$$

④ Hoeffding: $W_1, \dots, W_n$ independent & $a_i \leq W_i \leq b_i \quad \forall i$. Then $\forall t \geq 0 \quad Z = \frac{1}{n} \sum_{i=1}^n W_i$

$$P(|Z - E[Z]| \geq t) \leq 2 \exp\left(-\frac{2n^2 t^2}{\sum_{i=1}^n (b_i - a_i)^2}\right)$$

Finite Hypothesis classes: (Thm 7) Assume $\ell(h(x), y) \in [0, L], \quad \forall h \in \mathcal{H}$. $\bar{h} = \arg\min_{h \in \mathcal{H}} R(h)$
 $|\mathcal{H}| < \infty$. Then w/ prob. at least $1 - \delta$, the ERM $\hat{h}$ satisfies

$$R(\hat{h}) - R(\bar{h}) \leq L \sqrt{\frac{2 \log(|\mathcal{H}|) + \log(\delta^{-1})}{n}}$$

Pf: If $\hat{h} = \bar{h} \Rightarrow R(\hat{h}) - R(\bar{h}) = 0 \Rightarrow P(R(\hat{h}) - R(\bar{h}) > t) = P(R(\hat{h}) - R(\bar{h}) > t, \hat{h} \neq \bar{h})$

By (*) $P(R(\hat{h}) - R(\bar{h}) > t) \leq P(R(\hat{h}) - \hat{R}(\hat{h}) + \hat{R}(\bar{h}) - R(\bar{h}) \geq t, \hat{h} \neq \bar{h})$

$$\leq P(R(\hat{h}) - \hat{R}(\hat{h}) > t_2, \hat{h} \neq \bar{h}) + P(\hat{R}(\bar{h}) - R(\bar{h}) > t_2)$$

Bound (***) w/ Hoeffding. $\hat{R}(\bar{h}) - R(\bar{h}) = \frac{1}{n} \sum_{i=1}^n \ell(\bar{h}(x_i), y_i) - E[\ell(\bar{h}(x_i), y_i)]$

So

$$P(\hat{R}(\bar{h}) - R(\bar{h}) > t_2) \leq \exp\left(-\frac{2n^2 (t_2)^2}{n L^2}\right) = e^{-\frac{n t_2^2}{2 L^2}}$$

when $\hat{h} \neq \bar{h}$, $R(\hat{h}) - \hat{R}(\hat{h}) \leq \max_{h \in \mathcal{H} \setminus \{\bar{h}\}} R(h) - \hat{R}(h)$, so

$$(**) \leq P(\exists h \in \mathcal{H} \setminus \{\bar{h}\} : R(h) - \hat{R}(h) > t_2) \leq \sum_{h \in \mathcal{H} \setminus \{\bar{h}\}} P(R(h) - \hat{R}(h) > t_2) \quad [\text{union bd}]$$

$$\leq (|\mathcal{H}| - 1) \exp\left(-\frac{n t_2^2}{2 L^2}\right) \quad [\text{Hoeffding}]$$

$$\therefore P(R(\hat{h}) - R(\bar{h}) > t) \leq |\mathcal{H}| \exp\left(-\frac{n t^2}{2 L^2}\right) \leq \delta$$

$$\Leftrightarrow -\log\left(\frac{|\mathcal{H}|}{\delta}\right) \geq -\frac{n t^2}{2 L^2}$$

$$\Leftrightarrow t^2 \geq 2 L^2 \log(|\mathcal{H}|/\delta) / n$$

Note: $P(\forall h \in \mathcal{H}, R(h) - \hat{R}(h) \leq t) = 1 - P(\exists h \in \mathcal{H}, R(h) - \hat{R}(h) > t) \quad \delta = |\mathcal{H}| e^{-\frac{t^2 n}{2 L^2}}$

can use this idea in model selection (how to choose $\mathcal{H}$) w/ techniques known as structural risk minimization

$$\geq 1 - \sum_{h \in \mathcal{H}} P(R(h) - \hat{R}(h) > t) \quad \log\left(\frac{\delta}{|\mathcal{H}|}\right) = -\frac{t^2 n}{2 L^2}$$

$$\geq 1 - |\mathcal{H}| \exp\left(-\frac{t^2 n^2}{n L^2}\right) \quad [\text{Hoeffding}]$$

$$E[\ell(h(x), y)]$$

$$\therefore \forall h \in \mathcal{H}, \underbrace{P(R(h) < \hat{R}(h) + \sqrt{\frac{L^2 \log(|\mathcal{H}|/\delta)}{2n}})}_{\text{cannot be computed}} \geq 1 - \delta \quad \underbrace{\left(\text{evaluate w/ data}\right)}_{\text{evaluate w/ data}} \quad [\text{holds for } h = \hat{h}]$$

Example: $X = [0, 1]^2$ partitioned into m^2 disjoint squares width $\frac{1}{m}$ each. $|\mathcal{H}| = 2^{m^2}$ & apply

Infinite Hypothesis classes: $|\mathcal{H}| = \infty$, As $R(\hat{h}) - R(\bar{h}) \leq R(\hat{h}) - \hat{R}(\hat{h}) - \hat{R}(\bar{h}) - R(\bar{h})$

$$E[R(\hat{h}) - R(\bar{h})] \leq E[G] \quad G := \sup_{h \in \mathcal{H}} R(h) - \hat{R}(h)$$

take expectation

GOAL: bound $E[G]$ & hence bound the expected risk

$$Y = \{-1, 1\} \Rightarrow \text{BINARY CLASSIFICATION}$$

Set $z_i = (x_i, y_i)$, consider $\mathcal{F} = \{f(x, y) \mapsto \ell(h(x), y); h \in \mathcal{H}\}$, $G = \sup_{f \in \mathcal{F}} \frac{1}{n} \sum_{i=1}^n E[f(z_i)] - f(z_i)$

can write this for each ...

→ Rademacher Complexity: $\mathcal{F}$ class of $f: \mathcal{Z} \rightarrow \mathbb{R}$, $z_1, \dots, z_n \in \mathcal{Z}$

① $F(z_{1:n}) = \{(f(z_1), \dots, f(z_n)) : f \in \mathcal{F}\} \subset \mathbb{R}^n$ is a collection of vectors giving the behaviours of $\mathcal{F}$ on $z_{1:n}$.

② The empirical rademacher complexity is

$$\hat{\mathcal{R}}(\mathcal{F}(z_{1:n})) := \frac{1}{n} \mathbb{E} \left[\sup_{f \in \mathcal{F}} \sum_{i=1}^n \varepsilon_i f(z_i) \right]$$

"how closely aligned is your classifier on your dataset?"

w/ $(\varepsilon_i)_{i=1}^n$ iid rademacher. This quantifies how close $f(z_{1:n})$ is to random labels.

③ Take $z_1, \dots, z_n$ as R.V. The rademacher complexity is

$$\mathcal{R}(\mathcal{F}) := \mathbb{E}[\hat{\mathcal{R}}(\mathcal{F}(z_{1:n}))]$$

$$\hat{\mathcal{R}}(\mathcal{F}(z_{1:n})) = \frac{1}{n} \mathbb{E} \left[\sup_{f \in \mathcal{F}} \sum_{i=1}^n \varepsilon_i f(z_i) \mid z_{1:n} \right]$$

Some examples... pick sup f carefully based on situation. Idea: bd $\hat{\mathcal{R}}(\mathcal{F}(z_{1:n}))$ uniformly in $z_{1:n}$

→ Thm 8: In setting above, $\mathbb{E} \left[\sup_{f: \mathcal{Z} \rightarrow \mathbb{R}} \frac{1}{n} \sum_{i=1}^n (\mathbb{E}[f(z_i)] - f(z_i)) \right] \leq 2\mathcal{R}_n(\mathcal{F})$

depends on z_i distribution... complex

independent of P_0

Pf: immediate consequence

→ Thm 9: $\mathcal{F} := \{(x, y) \mapsto \varphi(h(x), y) : h \in \mathcal{H}\}$. Then

$$\mathbb{E}[R(\hat{h}) - R(\bar{h})] \leq 2\mathcal{R}_n(\mathcal{F})$$

MAIN RESULT OF SECTION

THE ESTIMATION ERROR IS CONTROLLED BY THE RAOEMACHER COMPLEXITY

Example 9 (Linear Models): If $\varphi(\cdot, y)$ is G -Lipschitz, then $\mathcal{F} = \{\varphi(h(x), y), h \in \mathcal{H}\}$ satisfies $\mathcal{R}_n(\mathcal{F}) \leq G\mathcal{R}_n(\mathcal{H})$. Take $\mathcal{H} = \{h_w(x) = w^T \varphi(x) : \|w\| \leq k\}$. Then ...

End up w/

$$\mathbb{E}[R(\hat{h}) - \inf_{h \in \mathcal{H}} R(h)] \leq \frac{k \sqrt{\mathbb{E}[\|\varphi(x)\|^2]}}{\sqrt{n}}$$

independent of φ dimension!

long & hard... didn't come up...

→ VC Dimension: Bounding expected risk → bounding rademacher complexity

study w/ VC dimension

Idea: given data $z_{1:n}$, count behaviours of $\mathcal{F}$: $|F(z_{1:n})|$

Note: $|F(z_{1:n})| \leq |\mathcal{H}(z_{1:n})|$

Lemma 11: $w_1, \dots, w_d$ mean zero s.g. $\Rightarrow \mathbb{E}[\max_j w_j] \leq \sigma \sqrt{2 \log(d)}$

Pf: use MGF, add more terms & optimise

Have an injection $\mathcal{F} \rightarrow \mathcal{H}$
 $(\varphi(h(x_i), y_i))_{i=1}^n \neq (\varphi(h'(x_i), y_i))_{i=1}^n$
 $\downarrow$
 $h(x_i) \neq h'(x_i)$

Lemma 10 (Massart's Lemma): $\mathcal{F} = \{(x, y) \mapsto \varphi(h(x), y), h \in \mathcal{H}\}$, φ takes values in $[0, 1]$

Then

$$\hat{\mathcal{R}}(\mathcal{F}(z_{1:n})) \leq \sqrt{\frac{2 \log |F(z_{1:n})|}{n}} \leq \sqrt{\frac{2 \log |\mathcal{H}(z_{1:n})|}{n}}$$

order $\frac{1}{\sqrt{n}}$ convergence just like the finite case for $|\mathcal{H}|$

Pf: By def have $\mathbb{E}[\sup_{f \in \mathcal{F}(z_{1:n})} \frac{1}{n} \sum_{i=1}^n \varepsilon_i f(z_i)]$, so need exactly $\mathbb{E}[\max_{j=1, \dots, d} w_j] \leq \sigma \sqrt{2 \log(d)}$

Let $d = |F(z_{1:n})|$ finite family

Let $\mathcal{F}' = \{f_1, \dots, f_d\}$ s.t. $\mathcal{F}'(z_{1:n}) = F(z_{1:n})$

So let $w_j = \frac{1}{n} \sum_{i=1}^n f_j(z_i) \varepsilon_i \Rightarrow \mathbb{E}[\max_{j=1, \dots, d} w_j] \leq \sqrt{\frac{2 \log(d)}{n}}$

Rademacher R.V. sub-gaussian w/ $\sigma=1$ so apply

NOTE: $\mathbb{E}[R(\hat{h}) - \inf_{h \in \mathcal{H}} R(h)] \leq 2\mathcal{R}_n(\mathcal{F}) \leq 2 \mathbb{E} \left[\sqrt{\frac{2 \log |\mathcal{H}(z_{1:n})|}{n}} \right]$

top grows slower than bottom to be useful

If $\gamma \in \{-1, 1\}^n$, $\mathcal{H}(z_{1:n}) = \{(h(x_i))_{i=1}^n : h \in \mathcal{H}\} \Rightarrow |\mathcal{H}(z_{1:n})| \leq 2^n$ **BAD!**

GOAL: characterise classes of functions where $|\mathcal{H}(z_{1:n})| \leq 2^n \Rightarrow$ Massart gives bound.

otherwise, estimation error unbd $\Rightarrow$ useless class of functions!

Def: Let $\mathcal{H}$ be a class of functions $h: X \rightarrow \{a, b\}$, $a \neq b$. $|\mathcal{H}| \geq 2$.

VC Dimension definition

- $\mathcal{H}$ shatters $x_{1:n} \in X^n$ if $|\mathcal{H}(x_{1:n})| = 2^n$ "every possible labelling is achieved on $x_{1:n}$ "
- $\Pi_{\mathcal{H}}(n) = \max_{x_{1:n} \in X^n} |\mathcal{H}(x_{1:n})|$ is the shattering coefficient / growth function "max # behaviours on n data points"
- The VC dimension of $\mathcal{H}$, $VC(\mathcal{H}) = \sup \{n \in \mathbb{N} : \Pi_{\mathcal{H}}(n) = 2^n\}$ "largest # data points s.t. $x_{1:n}$ is shattered"

largest # datapoints that achieve every possible labelling

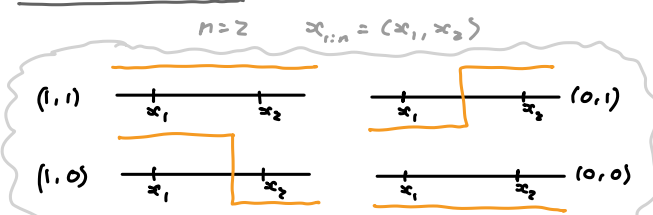
To check $VC(\mathcal{X}) = n$: ① Find some $x_{1:n}$ that can be shattered by $\mathcal{X}$ ② Show no set of $n+1$ pts can be shattered by $\mathcal{X}$

Example: $X = \mathbb{R}$, $\mathcal{Y} = \{0,1\}$, $\mathcal{X} = \{ \mathbb{I}_{[a,b)} : a, b \in \mathbb{R} \}$, Given data $x_1, \dots, x_n$, how many distinct vectors in $\mathcal{X}(x_{1:n})$? If a & a' in the same interval $(x_i, x_{i+1}] \Rightarrow \mathbb{I}_{[a,b)}(x_i) = \mathbb{I}_{[a',b)}(x_i)$. There are $n+1$ intervals $(-\infty, x_1], (x_1, x_2], \dots, (x_n, \infty)$. If b & b' in the " " $(x_j, x_{j+1}] \Rightarrow \mathbb{I}_{[a,b)}(x_j) = \mathbb{I}_{[a',b')} (x_j)$. So we get the same behaviour.

Note: To show VC dim

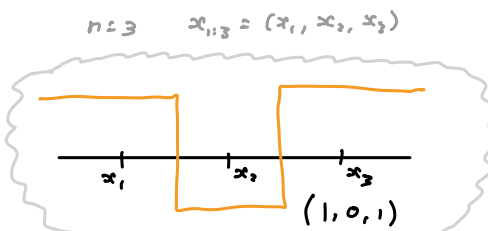
So as an upper bd, $n+1$ intervals to place a , $n+1$ intervals to place $b \Rightarrow |\mathcal{X}(x_{1:n})| \leq (n+1)^2$

VC DIMENSION: we set $|\mathcal{X}(x_{1:n})| \leq (n+1)^{VC(\mathcal{X})}$



$x_{1:2}$ can be shattered
 $\Rightarrow$ all possible labellings achieved

$\Rightarrow VC(\mathcal{X}) = 2$



$x_{1:3}$ cannot be shattered
 $\Rightarrow (1,0,1)$ labelling not achieved in $\mathcal{X}$

Pf: Non-exampl.

Sauer-Shelah lemma: If a class w/ finite VC dimension d , then $\Pi_{\mathcal{X}}(n) \leq (n+1)^d$

So, overall, $|\mathcal{X}(x_{1:n})| \leq \Pi_{\mathcal{X}}(n) \leq (n+1)^d$ where $d = VC(\mathcal{X})$, so

$$\underbrace{\mathbb{E}[R(\hat{h}) - \inf_{h \in \mathcal{H}} R(h)]}_{\text{expected estimation error}} \leq \underbrace{2\mathcal{R}(\mathcal{F})}_{\text{Minkowski's lemma}} \leq 2 \sqrt{\frac{2 VC(\mathcal{X}) \log(n+1)}{n}} \xrightarrow{n \rightarrow \infty} 0$$

CHAPTER TWO

Optimisation

$$\text{argmin}(f) = \{w \in \mathbb{R}^p : f(w) = \inf_{w \in \mathbb{R}^p} f(w)\}$$

⑧ Unconstrained Optimisation: $f: \mathbb{R}^p \rightarrow \mathbb{R}$, goal: $\inf_{w \in \mathbb{R}^p} f(w)$

$w \in \mathbb{R}^p$ is a

$\hookrightarrow$ global minimiser if $\forall w \in \mathbb{R}^p$ $f(w^*) \leq f(w)$

$\hookrightarrow$ local " " $\exists$ open neighbourhood U of w^* s.t. $f(w^*) \leq f(w) \quad \forall w \in U$

$\hookrightarrow$ strict local " " " " " " " " $f(w^*) < f(w) \quad \forall w \in U \setminus \{w^*\}$

$\hookrightarrow$ isolated local minimiser if " " " " " in which w^* is the only local minimiser

• A set C is convex if $\forall w, v \in C, \forall \lambda \in [0,1], \lambda v + (1-\lambda)w \in C$

• A function $f: S \rightarrow \mathbb{R}$ convex if $f(\lambda v + (1-\lambda)w) \leq \lambda f(v) + (1-\lambda)f(w) \quad \forall v, w \in S, \forall \lambda \in [0,1]$

$\hookrightarrow$ concave if $-f$ convex

• $f: \mathbb{R}^p \rightarrow \mathbb{R}$ convex. Then any local minimiser is a global minimiser. If strictly convex $\uparrow$

Pf: probs won't come up...

• Recall that $f: \mathbb{R}^p \rightarrow \mathbb{R}$ diff $\Rightarrow f(w+v) = f(w) + \langle \nabla f(w), v \rangle + o(\|v\|)$ [Taylor]

• Characterisation of convexity via differentiability:

(i) If f diff. f convex $\Leftrightarrow \forall w, v \in \mathbb{R}^p$ $f(w) \geq f(v) + \nabla f(v)^T (w-v)$

$$v^T \nabla f(v) \geq 0 \quad \forall v$$

(ii) If f twice diff f convex $\Leftrightarrow$ Hessian $\nabla^2 f(v)$ positive semi-definite $\forall v \in \mathbb{R}^p$

Pf: probs won't come up...

• First order optimality condition: f convex & diff. Then $w^* \in \text{argmin}(f) \Leftrightarrow \nabla f(w^*) = 0$

Pf: ' $\Rightarrow$ ' well known, ' $\Leftarrow$ ' convexity $\Rightarrow f(w) \geq f(w^*) + \nabla f(w^*)^T (w-w^*) = f(w^*) \Rightarrow w^*$ minimiser

Note: $f(x) = -x^2$. $\nabla f(0) = 0$ but 0 maximiser

$f(x) = x^3$ $\nabla f(0) = 0$ but 0 not local minimiser

$$\begin{aligned} (A-w-b)^T(A-w-b) &= (A-w-b)^T(A-w-b) \\ &= f(w) + \langle A-w-b, A \rangle + \langle A-w-b, -b \rangle = 2A^T(A-w-b) \end{aligned}$$

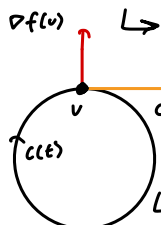
Example (Least Squares): $f(w) = \|Aw - b\|^2$ $A \in \mathbb{R}^{n \times p}$, $b \in \mathbb{R}^n$. As $f(x+\varepsilon) = f(x) + \nabla f(x)^T \varepsilon$

$$f(w+\varepsilon) = \|Aw - b + A\varepsilon\|^2 = \|Aw - b\|^2 + 2\langle Aw - b, A\varepsilon \rangle + \|A\varepsilon\|^2 \Rightarrow \nabla f(w) = 2A^T(Aw - b)$$

Convex $\Rightarrow w^*$ minimiser iff $A^T A w^* = A^T b$. If $A^T A$ invertible ($\text{ker}(A^T A) = \{0\}$),

Then $w^* = (A^T A)^{-1} A^T b$. [Note $\nabla^2 f = 2A^T A$ is useful for true diff]

• Gradient descent: can't compute $w^* \in \text{argmin}(f)$ in closed form $\Rightarrow$ construct $w_k \xrightarrow{k \rightarrow \infty} w^*$
 [inverting big matrices]



→ Direction of steepest descent: Require $\langle \nabla f(x_0), p \rangle < 0$ as $-\nabla f(x_0)$ is the direction of steepest descent. **NOTE**: $\nabla f(v)$ is orthogonal to level sets. GRADIENT DESCENT

Initialize w/ $w_0 \in \mathbb{R}^p$. For $k=0,1,2,3,\dots$ $w_{k+1} = w_k - \tau_k \nabla f(w_k)$ learning rate / step-size

→ Choosing the stepsize: Not too small [converge ASAP], not too big [Increase f, not converge]

(A) Greedy choice: $\tau = \arg \min h(\tau)$, $h(\tau) = f(w_k - \tau \nabla f(w_k))$

compute $h'(\tau)$ set $h'(\tau_k) = 0$

$$h(\tau + \epsilon) = f(w_k - \tau \nabla f(w_k) - \epsilon \nabla f(w_k))$$

$$= f(w_k - \tau \nabla f(w_k)) + \epsilon \nabla f(w_k)^T \nabla f(w_k - \tau \nabla f(w_k)) + o(\epsilon)$$

$$\Rightarrow h'(\tau) = \nabla f(w_k)^T \nabla f(w_k - \tau \nabla f(w_k)) \stackrel{!}{=} 0 \Rightarrow \nabla f(w_k)^T \nabla f(w_{k+1}) = 0$$

As $\nabla f(w_k) \nabla f(w_{k+1}) = 0 \Rightarrow$ orthogonal gradient descent directions $\Rightarrow$ zig-zag



(B) Armijo rule: Find τ that sufficiently decreases f. some backtracking line search.

Example (least squares): Greedy choice $\Rightarrow 0 = \nabla f(w_k)^T \nabla f(w_{k+1}) = \langle r_k, r_{k+1} \rangle$

$$w_{k+1} = w_k - \tau_k r_k, \quad r_k = \nabla f(w_k) = A^T(Aw_k - b) \Rightarrow \dots \Rightarrow \tau_k = \frac{\|r_k\|^2}{\|Ar_k\|^2}$$

→ Convergence Analysis for GD: [don't need proofs for exam, state results here]...

• If $0 < \tau_{\min} \leq \tau_k \leq \tau_{\max} < \frac{2}{L}$. Then $\exists \rho \in [0,1)$ s.t. $\|w_k - w^*\| \leq \rho^k \|w_0 - w^*\|$

[convex problem] $\Rightarrow$ ie w_k converges linearly to w^* [L largest eigenvalue for $C: f(w) := \frac{1}{2} \langle Cw, w \rangle - \langle w, b \rangle$]

• $f: S \rightarrow \mathbb{R}$ diff, $S \subset \mathbb{R}^p$ convex

(i) f is μ -strongly convex if $\forall x, x' \in S, \langle \nabla f(x) - \nabla f(x'), x - x' \rangle \geq \mu \|x - x'\|^2$

(ii) f is L -Lipschitz smooth if $\|\nabla f(x) - \nabla f(x')\| \leq L \|x - x'\|$

$\Rightarrow$ These imply f can be lower & upper bdd by quadratics...

Main result on convergence of GD

when f satisfies (i)&(ii) $\Rightarrow$ linear rate of convergence $\|w_k - w^*\| \leq \rho^k \|w_0 - w^*\|$

when f only (ii) (Lipschitz smooth) $\Rightarrow O(\frac{1}{k})$ convergence [slower] $f(w_k) - f(w^*) \leq \frac{C}{k}$

when f Lipschitz cts $\Rightarrow O(\frac{1}{\sqrt{k}})$ convergence [even worse]

• Stochastic Gradient Descent: what if for each iteration k , draw an index i uniformly at random from $\{1, \dots, n\}$? Problem: $\min_{w \in \mathbb{R}^p} f(w)$ where $f(w) = \frac{1}{n} \sum_{i=1}^n f_i(w)$

$$\mathbb{E}[\nabla f_{i_k}(w)] = \sum_{j=1}^n \nabla f_j(w) \cdot \mathbb{P}(i=j) = \frac{1}{n} \sum_{j=1}^n \nabla f_j(w) = \nabla f(w)$$

e.g. $\min_{w \in \mathbb{R}^p} \frac{1}{n} \sum_{i=1}^n \ell(w(x_i), y_i)$

$\ell(z, y) = \frac{1}{2} (z - y)^2$

$f_i(w) = \frac{1}{2} (w^T x_i - y_i)^2$

Highly relevant to ML

So, replace expensive gradient computation by $\nabla f_{i_k}(w)$

Stochastic Gradient Descent

Initialize $w_0 \in \mathbb{R}^p$, for $k=0,1,2,\dots$ $w_{k+1} = w_k - \nabla f_{i_k}(w_k)$ $i_k \stackrel{\text{unit}}{\sim} \{1, \dots, n\}$ $\{w_k\}$ random vectors

Example: $f_i(w) = \frac{1}{2} (w^T x_i - y_i)^2 \Rightarrow \nabla f_i(w) = (w^T x_i - y_i) x_i$. If $w \in \mathbb{R}^p$ $\{x_i\}_{i=1}^n$, n data points, then each k of SGD costs $O(p)$, each k of GD $\Rightarrow \nabla f(w) \Rightarrow O(np) \Rightarrow n$ times cheaper!

→ Thm 19: suppose f μ -strongly convex, $\|\nabla f_i(w)\|^2 \leq C^2 \quad \forall i$. Let $\tau_k = \frac{1}{\mu(k+1)}$

Then $\mathbb{E}[\|w_k - w^*\|^2] \leq \frac{R}{k+1}$ where $R = \max \left\{ \|w_0 - w^*\|^2, \frac{C^2}{\mu^2} \right\}$

Remark:

- SGD much slower than GD. Only weakly benefits from strong convexity & no linear rate
- Advantage: SGD makes progress towards w^* w/out full pass through n data points. (w/ large datasets, you might only make 1 full pass).

(P) Constrained Optimisation: Consider $\min_{w \in \mathbb{R}^p} f_0(w)$ subject to $Aw = b$ and $\forall i=1,\dots,m, f_i(w) \leq 0$ w/ $A \in \mathbb{R}^{n \times p}$, $b \in \mathbb{R}^n$, $f_i: \mathbb{R}^p \rightarrow \mathbb{R}$. Assume f_i convex

• Feasibility set of (P) $\mathcal{F} := \{w \in \mathbb{R}^p: \forall i \in \{1, \dots, m\}, f_i(w) \leq 0 \text{ and } Aw = b\}$ place where all the conditions hold

• First order optimality condition: Let f be convex & diff. consider $\min_{w \in \mathcal{F}} f(w)$ for $\mathcal{F}$ convex.

Then w^* minimiser $\Leftrightarrow \forall w \in \mathcal{F}, \nabla f(w^*)^T (w - w^*) \geq 0$

Pf: $f(w) \geq f(w^*) + \langle \nabla f(w^*), w - w^* \rangle \geq f(w^*) \Rightarrow w^*$ minimiser
 $\Rightarrow$ contradiction. Assume w^* minimiser but $\nabla f(w^*)^T (w - w^*) < 0$ for some $w \in \mathcal{F}$
 Let $\lambda \in [0, 1]$, define $v(\lambda) = (1-\lambda)w^* + \lambda w \in \mathcal{F}$. Define $g(\lambda) = f(v(\lambda))$
 $g'(\lambda) = \nabla f(v(\lambda))^T (w - w^*)$ [chain rule] $\Rightarrow g'(0) = \langle \nabla f(w^*), w - w^* \rangle < 0$ [assumption]
 so g decreasing at $\lambda=0 \Rightarrow f(v(\lambda)) < f(v(0)) = f(w^*)$ $\forall \lambda$ sufficiently small. $\times$

$\hookrightarrow$ Duality: Primal (P) & dual offer different perspectives to help solve

The Lagrange function of (P) is defined $\forall w \in \mathbb{R}^p, \xi \in \mathbb{R}^n, \nu \in \mathbb{R}_{\geq 0}^m$ by

$$L(w, \xi, \nu) = f_0(w) + \underbrace{\xi^T (Ax - b)}_{\text{linear}} + \sum_{k=1}^m \underbrace{\nu_k f_k(w)}_{\text{inequalities}}$$

ξ, ν are the Lagrange multipliers

The Lagrange dual function is $D(\xi, \nu) = \inf_{w \in \mathbb{R}^p} L(w, \xi, \nu)$

note: if $w \mapsto L(w, \xi, \nu)$ unbounded from below
 $D(\xi, \nu) = -\infty$

$\hookrightarrow$ Property 1: Dual function is concave: $(\xi, \nu) \mapsto D(\xi, \nu)$ linear in $\xi, \nu \Rightarrow$ ptwise infimum of linear affines is concave.

$\hookrightarrow$ Property 2: Dual lower bounds $\inf_{w \in \mathcal{F}} f_0(w)$

$$D(\xi, \nu) \leq L(w, \xi, \nu) = f_0(w) + \underbrace{\xi^T (Ax - b)}_{=0} + \sum_{k=1}^m \underbrace{\nu_k f_k(w)}_{\leq 0} \leq f_0(w) \quad \forall w \in \mathcal{F}$$

$$\text{Taking supremum} \Rightarrow D(\xi, \nu) \leq \inf_{w \in \mathbb{R}^p} f_0(w)$$

The dual problem to (P) is

$$\max_{\xi \in \mathbb{R}^n, \nu \in \mathbb{R}^m} D(\xi, \nu) \quad \text{s.t. } \nu \geq 0 \quad (D)$$

Let (ξ^*, ν^*) maximise (D) and w^* minimise (P)

$\hookrightarrow$ weak duality: $D(\xi^*, \nu^*) \leq f_0(w^*)$

$\hookrightarrow$ strong duality: $D(\xi^*, \nu^*) = f_0(w^*)$

KKT CONDITIONS

- Feasibility conditions
- inequalities $\Rightarrow$ w/ multipliers
- $\nabla L = 0$

when does strong duality hold?

Slater's constraint qualification: Assume f_i convex $\forall i \in \{0, \dots, m\}$. $\text{dom}(f_0) = \mathbb{R}^p$. If $\exists w \in \mathbb{R}^p$ s.t.

$Aw = b$ and $f_k(w) < 0 \quad \forall k \in \{1, \dots, m\}$. Then strong duality holds $\exists$ pt in the INTERIOR

If we have strong duality, how are ξ^*, ν^*, w^* related?

If w^*, ν^*, ξ^* optimise (P) and (D) & strong duality holds, then following KKT hold:

$$f_i(w^*) \leq 0, Aw^* = b, \nu_i^* \geq 0$$

$$\nu_i^* f_i(w^*) = 0 \quad \forall i \in \{1, \dots, m\}$$

$$\nabla f_0(w^*) + \sum_i \nu_i^* \nabla f_i(w^*) + A^T \xi^* = 0$$

$$\text{why? } f_0(w^*) = D(\xi^*, \nu^*) = \inf_{w \in \mathcal{F}} f_0(w) + \underbrace{\langle \xi^*, Aw - b \rangle}_{=0 \quad \forall w \in \mathcal{F}} + \sum_{i=1}^m \underbrace{\nu_i^* f_i(w)}_{\nu_i^* \geq 0, f_i \leq 0} \leq f_0(w^*) + \sum_{i=1}^m \nu_i^* f_i(w^*) \leq f_0(w^*)$$

so $\leq \Rightarrow =$

yields

$$\sum_{i=1}^m \nu_i^* f_i(w^*) = 0 \Rightarrow \nu_i^* f_i(w^*) = 0 \quad \forall i \in \{1, \dots, m\} \quad [\text{complementary slackness}]$$

$$w \in \arg\min_w L(w, \xi^*, \nu^*) \Rightarrow \nabla_w L(w^*, \xi^*, \nu^*) = 0$$

(10) Support Vector Machines: Data $(x_i, y_i) \quad y_i \in \{-1, 1\}, x_i \in \mathbb{R}^p$. Goal:

How to pick h ? maximise distance to closest data point in both classes. Margin between $H, \{x_i\}$

$$\gamma(w, b) = \min_{1 \leq i \leq n} \{ \|x - x_i\|, i=1, \dots, n, x \in H \}$$

Distance to hyperplane. Given $x \in \mathbb{R}^p$, $\min_{z \in H} \|x - z\| = ?$

(A) Geometric argument: w orthogonal to hyperplane

so projection of x onto H has form $z = x + tw$ for some $t \in \mathbb{R}$.

$$\text{Require } z \in H \Rightarrow -b = w^T z = w^T x + t \|w\|^2$$

$$\Leftrightarrow t = \frac{-b - x^T w}{\|w\|^2}$$

$$\Rightarrow \min_{z \in H} \|x - z\| = |t| \|w\| = \frac{|b + x^T w|}{\|w\|}$$

linear classifier

$$h(x) = \text{sgn}(w^T x + b) \quad w \in \mathbb{R}^p, b \in \mathbb{R}$$

Note: $h(x_i) = y_i$

$$\Leftrightarrow \text{sgn}(w^T x_i + b) = y_i$$

$$\Leftrightarrow y_i (w^T x_i + b) > 0$$

$$y_i = -1, h(x_i) = -1 \rightarrow +ve$$

$$y_i = 1, h(x_i) = 1 \rightarrow +ve$$

so SVM is

$$\max_{w, b} \min_{1 \leq i \leq n} \frac{|b + x_i^T w|}{\|w\|}$$

$$\text{s.t. } y_i (w^T x_i + b) \geq 0$$

③ Dual Approach: Optimisation problem: $\min \|x - z\|^2$ s.t. $w^T z + b = 0$ [x is given data]

Lagrange function: $L(z, \xi) = \|x - z\|^2 + \xi(w^T z + b)$

Lagrange Dual function: $D(\xi) = \min_z \|x - z\|^2 + \xi(w^T z + b)$ $\frac{1}{4}\xi^2\|w\|^2 - \frac{1}{2}\xi^2\|w\|^2 + \xi(w^T x + b)$

Strong duality: Holds $\therefore$ convex problem & $\exists z$ s.t. $w^T z + b = 0$

First order optimality: Minimiser z^* of $D(\xi)$ satisfies $\nabla D(z^*) + A^T \xi = 0$

$$\text{So } 2(z^* - x) + \xi w = 0 \Rightarrow z^* = -\frac{1}{2}\xi w + x$$

Sub min into $D(\xi)$: $D(\xi) = -\frac{1}{4}\xi^2\|w\|^2 + \xi(w^T x + b)$

Dual problem $\max_{\xi \geq 0} D(\xi)$: maximiser ξ^* satisfies $\nabla D(\xi^*) = -\frac{1}{2}\xi^*\|w\|^2 + w^T x + b = 0$
 So $\xi^* = \frac{2(b + w^T x)}{\|w\|^2}$ $\therefore z^* = x - \frac{b + w^T x}{\|w\|^2} w$ [First order optimality]

Solve w/ optimal z^* : $\|z^* - x\|^2 = \min_{z \in H} \|z - x\|^2 = \frac{|b + w^T x|}{\|w\|}$

Maximum margin classifier for SVM:

Find w, b to max $= \gamma(w, b)$

$$\max_{w, b} \min_{1 \leq i \leq n} \frac{|b + w^T x_i|}{\|w\|}$$

s.t. $\forall i: x_i(w^T x_i + b) \geq 0$

See complexity

Scale invariance: $\gamma(tw, tb) = \gamma(w, b) \quad \forall t \in \mathbb{R} \setminus \{0\} \Rightarrow$ fix $\min_{1 \leq i \leq n} |b + w^T x_i| = 1$
 consider $\arg \max_{w, b} \frac{1}{\|w\|}$ s.t. $x_i(w^T x_i + b) \geq 0$ and $|b + w^T x_i| > 1$

Primal formulation: Absorb conditions: $\arg \min_{w, b} \|w\|$ s.t. $x_i(w^T x_i + b) \geq 1 \quad \forall i$ (P)

Dual Problem SVM: (P) is convex. Accurately separating H ex-sets $\Rightarrow$ HAVE STRONG DUALITY! Support vector machine. Data s.t. $x_i(w^T x_i + b) = 1$ support vectors

$\hookrightarrow$ Lagrange function: $L(w, b, \xi) = \frac{1}{2}\|w\|^2 + \sum \xi_i(1 - x_i(w^T x_i + b))$ for $\xi \in \mathbb{R}_{\geq 0}^n, (w, b) \in \mathbb{R}^{p+1}$

$\hookrightarrow$ Lagrange Dual: $D(\xi) = \inf_w L(w, b, \xi) = \inf_w \left(\frac{1}{2}\|w\|^2 + \sum \xi_i(1 - x_i(w^T x_i + b)) \right)$

$\hookrightarrow$ First order optimality: w, b minimise $D(\xi)$, then

$$\partial_w L(w, b, \xi) = w - \sum \xi_i x_i = 0, \quad \partial_b L(w, b, \xi) = -\sum \xi_i = 0$$

$\hookrightarrow$ Form of Dual:

$$D(\xi) = \begin{cases} -\frac{1}{2} \left\| \sum \xi_i x_i \right\|^2 + \sum \xi_i & \text{if } \sum \xi_i = 0 \\ -\infty & \text{orw} \end{cases}$$

$\hookrightarrow$ Then use geom of dual & complementary slackness: $\xi_i(1 - x_i(w^T x_i + b)) = 0$

Non-exact separation: see Q 6.3. You allow some mistakes.

Non-linear Decision boundaries: Project into higher dim space where you can have a linear separator. Define a feature map $\Phi: X \rightarrow \mathbb{R}^d$

$$\Phi(x) = (\phi_j(x))_{j=1}^d \quad \phi_j: X \rightarrow \mathbb{R} \text{ basis functions}$$

Hypothesis class:

$$\mathcal{H} = \{h(x) = w^T \Phi(x) + b : w \in \mathbb{R}^d, b \in \mathbb{R}\}$$

SVM in this setting becomes

$$\arg \min_{w \in \mathbb{R}^d, b \in \mathbb{R}} \|w\| \quad \text{s.t.} \quad x_i(w^T \Phi(x_i) + b) \geq 1 \quad (P)$$

$$\max_{\xi \in \mathbb{R}^n} -\frac{1}{2} \left\| \sum \xi_i \Phi(x_i) \right\|^2 + \sum \xi_i \quad \text{s.t.} \quad \sum \xi_i = 0 \text{ and } \xi_i \geq 0 \quad (D)$$

Point: dual problem reduces to evaluating the kernel. can avoid handling high dimensional Φ explicitly

Kernels: Goal is to avoid a high dim vector.

Define $k(x, x') = \Phi(x)^T \Phi(x')$

$$(D) \quad \left\| \sum \xi_i \Phi(x_i) \right\|^2 = \sum_{i,j} \xi_i \xi_j x_i^T x_j k(x_i, x_j)$$

$$= \langle \xi \circ Y, K_n(\xi \circ Y) \rangle$$

$$K_n = (k(x_i, x_j))_{i,j=1}^n \text{ [matrix]}, \quad \xi \circ Y = (\xi_i x_i)_{i=1}^n \text{ [vector]}$$

Examples:

① Polynomial kernels: $k(x, x') = x^T x'$ [linear], $k(x, x') = (x^T x' + 1)^d$ $d \in \mathbb{N}$ (polynomial)

E.g. $p=2, d=2, k(x, z) = (x^T z + 1)^2 = \left(\sum_{i=1}^2 x_i z_i + 1 \right)^2 = \sum_{i,j=1}^2 (x_i x_j)(z_i z_j) + \sum_{i=1}^2 (\sqrt{2} x_i)(\sqrt{2} z_i) + 1$

so feature vector: $\phi(x) = (x_1^2, x_1 x_2, x_2 x_1, x_2^2, \sqrt{2} x_1, \sqrt{2} x_2, 1)$

Point: $x \in \mathbb{R}^p, k(x, z) = (x^T z + 1)^d \rightarrow$ feature space has dimension $\binom{p+d}{d}$ [all monomials $\rightarrow$ deg=d]

while nothing in $O(p^d)$ dimensional space, computing $k(x, z)$ is $O(p)$

③ Gaussian kernel: $k(x, x') = \exp(-\frac{\|x - x'\|^2}{2\sigma^2})$ $\sigma > 0$

Practical considerations:

- You choose kernel. Not always obvious ①, ③ popular.
- choose d in ① / σ in ③ via cross validation. Big feature map $\Rightarrow$ high est. error low approx. error
- kernel SVM works best w/ small/medium sized datasets.

CHAPTER THREE

Neural Networks

see Qs for examples

applied pointwise

- ⑪ Multi-layer Perceptrons: consider functions of the form $h_{w,a,b}(x) = a^T \sigma(Wx + b)$
- Universal Approximation: Assume $\sigma: \mathbb{R} \rightarrow \mathbb{R}$ non-decreasing, cts. $\lim_{r \rightarrow -\infty} \sigma(r) = 0, \lim_{r \rightarrow \infty} \sigma(r) = 1$
- $\hookrightarrow$ **THM** For any compact set $\Omega \subset \mathbb{R}^p$, the space spanned by the functions $[w/\sigma \text{ as } \mathcal{J}]$

THE
UNIVERSAL
APPROXIMATION
THEOREM

$$\mathcal{F}_\sigma = \{ \varphi_{w,b}(x) := \sigma(x^T w + b), w \in \mathbb{R}^p, b \in \mathbb{R} \}$$

is dense in $C(\Omega)$ w/ uniform convergence. This means for ANY continuous function $f: \Omega \rightarrow \mathbb{R}$ and ANY $\varepsilon > 0$, $\exists q \in \mathbb{N}, w_j \in \mathbb{R}^p, a_j, b_j \in \mathbb{R}$ s.t.

$$\left| f(x) - \sum_{k=1}^q a_k \varphi_{w_k, b_k}(x) \right| \leq \varepsilon$$

Proof: exact proof non-examb. need the stone-weierstrass & how it can be applied...
[Fund. result in approximation theory]

- $\hookrightarrow$ The Stone Weierstrass Thm: Let $\Omega \subset \mathbb{R}^p$ be compact. Let $\mathcal{F}$ be a sub-algebra of $C(\Omega)$
- (i) $1 \in \mathcal{F}$, (ii) $f, g \in \mathcal{F} \Rightarrow \alpha f + \beta g \in \mathcal{F} \forall \alpha, \beta \in \mathbb{R}$ and $\mathcal{F}$ separates points in Ω . i.e. $\forall x \neq y \exists f \in \mathcal{F}$ s.t. $f(x) \neq f(y)$. Then $\mathcal{F}$ is dense in $C(\Omega)$ w/ uniform norm.
- i.e. $\forall g: \Omega \rightarrow \mathbb{R}$ cts, $\exists f \in \mathcal{F}$ s.t. $\|f - g\|_\infty \leq \varepsilon$ **Pf**: Fundamental result 1885

note: point separation necessary. If x, x' cannot be separated $\Rightarrow \forall f \in \mathcal{F} f(x) = f(x')$
Then define g s.t. $g(x) \neq g(x')$ & no f can approx.

E.g. $g(z) = \langle z - x, x' - x \rangle$. $g(x) = 0, g(x') = \|x' - x\|^2 \neq 0$

Example: $\forall g \in C(\Omega), \forall \varepsilon > 0, \exists f \in \text{Span } \mathcal{F}_{\exp}$ s.t. $\|f - g\|_\infty \leq \varepsilon$

Pf: verify stone-weierstrass conditions

- clearly all $f \in \mathcal{F}_{\exp}$ are continuous
- $x \mapsto \exp(\langle 0, x \rangle) = 1 \in \mathcal{F}_{\exp}$
- Span $\mathcal{F}_{\exp}$ is a linear subspace, hence an algebra: $e^{t+s} = e^t e^s$
- Point separation: if $x \neq y$, pick $f(z) = \exp(\langle z - x, y - x \rangle / \|y - x\|^2)$
 $f(x) = \exp(0) = 1, f(y) = \exp(1) = e$

- Deep Neural Networks: Idea is to stack several layers of perceptrons.

$F^l(x) = \sigma(W^l x + b^l), F^{k+1} = \sigma(W^{k+1} F^k(x) + b^{k+1}), W^k \in \mathbb{R}^{d_k \times d_{k-1}}, b^k \in \mathbb{R}^{d_k}$
call $F^l(x)$ a neural net of format $d = (d_1, \dots, d_L)$. $\mathcal{N} = \{ F: \mathbb{R}^{d_0} \rightarrow \mathbb{R}^{d_L} : F \text{ NN w/ format } d \}$

$\hookrightarrow$ classification: $d_L = K$ x assigned to argmax $F_i(x)$

$\hookrightarrow$ Regression: output $h(x) = a^T F^L(x)$ $a \in \mathbb{R}^{d_L}$

- ⑫ Automatic Differentiation: To find params of NN, $\min_{h \in \mathcal{H}} \frac{1}{N} \sum_{i=1}^N \ell(h(x_i), y_i)$ w/ SGD. Gradient bottleneck.

- Finite difference only gives an approx, symbolic diff memory intensive. Auto diff is exact. complexity same as eval. function
- Computational graphs
- Forward mode
- Reverse mode

Example: $f(z_1, z_2) = (\sin(z_1, z_2), \log(z_1) + z_1, z_2)$ $\frac{\partial f}{\partial z_k} = \sum_{\text{method}(h)} \frac{\partial f}{\partial z_m} \frac{\partial z_m}{\partial z_k} = \sum$

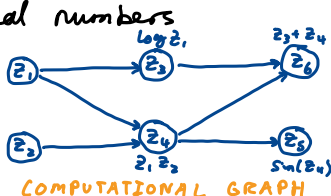
FORWARD PASS $f(2,3)$

$$\begin{aligned} z_1 &= 2, \\ z_2 &= 3 \\ z_3 &= \log 2 \\ z_4 &= 6 \\ z_5 &= \sin(6) \\ z_6 &= \log 2 + 6 \end{aligned}$$

BACKWARD PASS: Initialize

$$\begin{aligned} \dot{z}_6 &= 1, \dot{z}_5 = 1 \\ \dot{z}_4 &= \dot{z}_6 \cdot \frac{\partial f_6}{\partial z_4} + \dot{z}_5 \cdot \frac{\partial f_5}{\partial z_4} = 1 \times 1 + 1 \times \cos(z_4) = 1 + \cos(6) \\ \dot{z}_3 &= \dot{z}_6 \cdot \frac{\partial f_6}{\partial z_3} = 1 \times 1 = 1 \\ \dot{z}_2 &= \dot{z}_4 \cdot \frac{\partial f_4}{\partial z_2} = (1 + \cos(6)) \cdot z_1 = 2 + 2\cos(6) \\ \dot{z}_1 &= \dot{z}_3 \cdot \frac{\partial f_3}{\partial z_1} + \dot{z}_4 \cdot \frac{\partial f_4}{\partial z_1} = 1 \times \frac{1}{z_1} + (1 + \cos(6)) \times z_2 = \frac{1}{2} + 3 + 3\cos(6) \end{aligned}$$

So $\nabla f(2,3) \begin{pmatrix} 1 \\ 1 \end{pmatrix} = \begin{pmatrix} \frac{3\cos 6}{\frac{1}{2}+3} & 2\cos 6 \end{pmatrix} \begin{pmatrix} 1 \\ 1 \end{pmatrix} = \begin{pmatrix} 5\cos 6 \\ 5 + \frac{1}{2} \end{pmatrix}$



→ Gradient descent is a special case of this w/ $f: \mathbb{R}^p \rightarrow \mathbb{R}$.

- (13) Cross-Entropy Loss: Two prob. distributions P, Q on some sample space Ω . (discrete)
- KL-divergence between P & Q is $KL(P||Q) = \sum_{w \in \Omega} \log\left(\frac{P(w)}{Q(w)}\right) P(w)$
 - $Q(w) = 0 \Rightarrow P(w) = 0$
 - $KL(P||Q) \geq 0$ and $KL(P||Q) = 0 \iff P = Q$
 - k class classification:
 - To output a probability vector, post process via $p_k = \frac{e^{v_k}}{\sum_{j=1}^k e^{v_j}}$ where $F(x) = (v_1, \dots, v_k)$
 - $\ell(F(x), y) = - \sum_{j=1}^k y_j \log(p_j)$ [cross entropy loss]
- note: Binary classification $\ell(F(x), y) = -y \log p - (1-y) \log(1-p)$, $p = \frac{1}{1 + e^{-v}}$
 minimize $\sum \log(1 + e^{-(y_i w^T x_i)})$

- (14) Convolutional Neural Networks: Use in image processing applications
- Convolutions / Fourier transforms
 - $(f * g)(t) = \int f(x)g(t-x)dx$
 - won't come up surely...
 - CNN: a feedforward NN $x^{k+1} = \sigma(w^k x^k + b^k)$, $x^k \in \mathbb{R}^{n_k \times d_k}$
 - where $n_k = \#$ spatial positions, $d_k = \#$ channels e.g. RGB $\Rightarrow d_k = 3$.
 - spatial position $\uparrow$ & to ensure translation invariance, linear operator is a convolution operator. [convolutions are the only linear operators that are translation invariant]
 - max pool to $\downarrow$ # pixels. [summary of info in a window of size r]
-

- (15) Generative Adversarial Networks: Goal: learn the distribution P_X of the input X
- GANs. You have
 - Data space X w/ pdf P_X
 - latent space Z w/ pdf P_Z [compressed rep. of P_X]
 - $G: Z \rightarrow X$ generator
 - $D: X \rightarrow [0,1]$ discriminator.
- GAN is a (D, G) pair
- GOALS:
- Discriminator D aims to discriminate data sampled from P_X and P_Z .
 - Generator G aims to generate data in X that's indistinguishable from P_X .
- pushforward measure $P_G = G_* P_Z$

$$\int_X f(x) P_G(x) dx = \int_Z f(G(z)) P_Z(z) dz \quad \forall f: X \rightarrow \mathbb{R} \text{ measurable.}$$

Aim: $P_G = P_X$.

$$\min_G \max_D V(G, D) \quad , \quad V(G, D) = \int \log(D(x)) P_X(x) dx + \int \log(1 - D(x)) P_G(x) dx$$

⑥ Variational Auto-Encoders [compact representation of high dimensional data]

- why? To visualize data, ↓ noise, ↓ resources,
↑ performance of ↓ stream factor (e.g. classification)

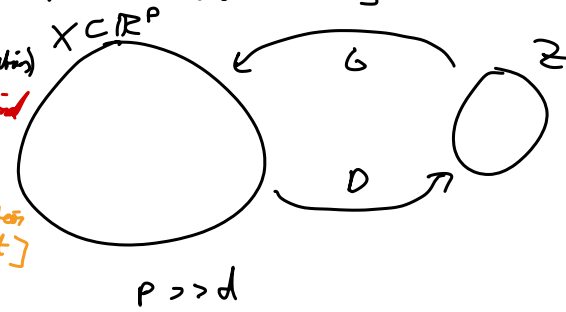
- Idea: Encoder $E: X \rightarrow Z$

$$\mathbb{R}^p \ni x \mapsto G(x) = z \in \mathbb{R}^d$$

$$\text{Decoder } D: Z \rightarrow X$$

$$z \mapsto P(z)$$

[Takes a representation & reconstructs it]



- Optimisation problem:

$$\min_{E, D} \frac{1}{n} \sum_{i=1}^n \|D(E(x_i)) - x_i\|^2 \quad \text{w/ data } \{x_i\}_{i=1}^n \subset \mathbb{R}^p$$

- VAE: model D & E as NNs,

$$\hookrightarrow z \sim q_\phi(\cdot | x) \quad \text{GIVEN Data}$$

$$\hookrightarrow x \sim p_\theta(\cdot | z) \quad \text{reconstruction}$$

[sample from these conditional probability distributions]

To generate data, new sample $z \sim p_z$ then draw from data space $p_\theta(\cdot | z)$

↳ Model the joint distribution as $p(x, z) = p_\theta(x | z) p_z(z)$

$$\hookrightarrow \text{Pick } p_z = \mathcal{N}(0, \mathbb{I}_d)$$

$$\hookrightarrow \text{model } p_\theta(\cdot | z) = \mathcal{N}(G_\theta(z), \sigma_\theta^2 \mathbb{I}_n) \quad \text{where } G_\theta: \mathbb{R}^d \rightarrow \mathbb{R}^p$$

This is good: p_z & $p(\cdot | z)$ simple gaussian

& can model complex data

[neural network]

$$p_\theta = \int p_\theta(x | z) p_z(z) dz$$

Useful Formulae

① $\mathbb{E}[x] = \mathbb{E}[\mathbb{E}[x|w]]$ [Tower property]

② $\mathbb{E}[g(w)z|w] = g(w)\mathbb{E}[z|w]$ [Taking out known]

③ $P(A|B) = \frac{P(B|A)P(A)}{P(B)}$ [Bayes Theorem] PF: $P(A|B) = \frac{P(A \cap B)}{P(B)}$, $P(B|A) = \frac{P(A \cap B)}{P(A)}$

④ $\inf_{\|w\| \leq \kappa} \{\|w - w_*\|\} = \max\{\|w_*\| - \kappa, 0\}$ [l.s. approximation error e.g.]

⑤ p is a descent direction of f if $\langle p, \nabla f \rangle < 0$

⑥ $f(x+\varepsilon) = f(x) + \langle \nabla f(x), \varepsilon \rangle + \text{h.o.t.}$ [Taylor expansion]

⑦ $\sigma(at) = \frac{e^{at}}{1+e^{at}} \xrightarrow{a \rightarrow \infty} \mathbb{I}_{[0,\infty)}$ [sigmoid is cts approx to heaviside]