

Week 1 Tues

quantify faithfulness / closeness

average each individual data point

Define loss function

$$\text{loss}(\bar{x}; x_i) = (\bar{x} - x_i)^2 \quad \leftarrow \text{squared loss / error}$$

Average loss function

$$R_{sq}(\bar{x}; D) = \frac{1}{n} \sum_{i=1}^n \text{loss}(\bar{x}; x_i) = \frac{1}{n} \sum_{i=1}^n (\bar{x} - x_i)^2$$

↑
empirical risk

Find an estimator h to optimize empirical risk?

arg min : argument x that minimize $f(x)$
 $f(h^*) \leq f(h), \forall h \in \mathbb{R}$

$$\bar{x} = \arg \min_{h \in \mathbb{R}} R_{sq}(h; D)$$

$$= \arg \min_{h \in \mathbb{R}} \frac{1}{n} \sum_{i=1}^n (h - x_i)^2$$

empirical risk minimization

mean!

$$h^* = \frac{1}{n} \sum_{i=1}^n x_i = \text{Mean}(D)$$

The best prediction, in terms of mean squared error, is the mean (unique)

more points example

$$2 \text{ point example: } \frac{\partial}{\partial h} \frac{1}{2} [(h - x_1)^2 + (h - x_2)^2]$$

$$= \frac{1}{2} \frac{\partial}{\partial h} ((h - x_1)^2 + (h - x_2)^2)$$

$$= \frac{1}{2} \left(\frac{\partial}{\partial h} (h - x_1)^2 + \frac{\partial}{\partial h} (h - x_2)^2 \right)$$

$$= \frac{1}{2} (2(h - x_1) + 2(h - x_2))$$

$$= \frac{1}{2} (2(h - x_1 + h - x_2))$$

$$= 2h - x_1 - x_2$$

↓

$$2h - x_1 - x_2 = 0$$

$$h = \frac{x_1 + x_2}{2}$$

$$\frac{\partial}{\partial h} \left(\frac{1}{n} \sum_{i=1}^n (h - x_i)^2 \right)$$

$$= \frac{1}{n} \left(\frac{\partial}{\partial h} \sum_{i=1}^n (h - x_i)^2 \right)$$

$$= \frac{1}{n} \sum_{i=1}^n \frac{\partial}{\partial h} (h - x_i)^2$$

$$= \frac{1}{n} \sum_{i=1}^n 2(h - x_i)$$

$$= \frac{2}{n} \sum_{i=1}^n (h - x_i)$$

$$= \frac{2}{n} \left(nh - \sum_{i=1}^n x_i \right)$$

||

$$\frac{2}{n} \left(nh - \sum_{i=1}^n x_i \right) = 0$$

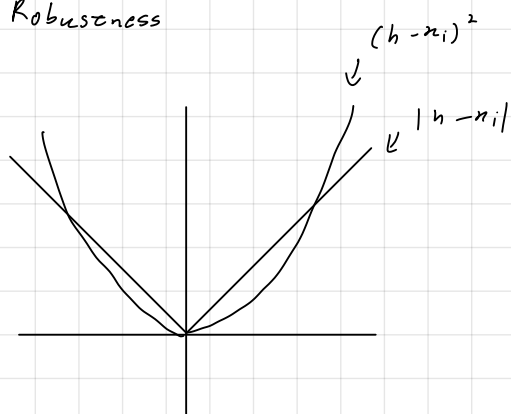
$$nh - \sum_{i=1}^n x_i = 0$$

$$\bar{x} = h^* = \frac{1}{n} \sum_{i=1}^n x_i$$

Outlier

Mean is sensitive to outliers.

Robustness



$$R^{ab}(h; D) = \frac{1}{n} \sum_{i=1}^n |h - x_i|$$

mean error
(good for robustness)

Week 1 Thurs

For dataset $D = \{x_1, \dots, x_n\}$

$$\text{loss}(h, x_i)$$

Empirical risk: $R(h, D) = \frac{1}{n} \sum_{i=1}^n \text{loss}(h, x_i)$

Optimize: $\underset{h \in \mathbb{R}}{\text{argmin}} R(h, D) = h^*$

training

Test: $R(h^*, D_{\text{test}})$

Def. A loss function $L(h, x)$ takes in a prediction h and a true value, x , and outputs a number measuring how far h is from x .

loss function:

- Squared loss - mean(D) - $\text{Loss}_{sq}(h, x_i) = (h - x_i)^2$

- Sensitive to outliers.

- Absolute loss: $\text{loss}_{abs}(h, x_i) = |h - x_i|$

↳ - Median(D)

optimize: $\underset{h \in \mathbb{R}}{\text{argmin}} R_{abs}(h, D) = \frac{1}{n} \sum_{i=1}^n \text{loss}_{abs}(h, x_i)$

$\lim_{\Delta \rightarrow 0} R_{abs}(h + \Delta, D) = R_{abs}(h, D)$

$$\frac{\partial}{\partial h} R_{abs}(h, D) = \frac{1}{n} \sum_{i=1}^n \frac{\partial}{\partial h} \text{loss}(h, x_i)$$

⇓

$$\mathbb{1}_{\{h < x_i\}} = \begin{cases} 1 & h < x_i \\ 0 & h > x_i \end{cases}$$

↓

$$\frac{\partial}{\partial h} \text{loss}(h, x_i) = \frac{\partial}{\partial h} |h - x_i| = \begin{cases} 1 & h > x_i \\ 0 & h = x_i \\ -1 & h < x_i \end{cases} \quad h \neq x_i$$

$$\approx -1 \cdot \mathbb{1}\{h < x_i\} + 1 \cdot \mathbb{1}\{h > x_i\}$$

$$\Rightarrow \frac{\partial}{\partial h} R_{\text{abs}}(h, D) = \frac{1}{n} \sum_{i=1}^n (-1 \cdot \mathbb{1}\{h < x_i\} + 1 \cdot \mathbb{1}\{h > x_i\})$$

$$= \frac{1}{n} \left[-1 \sum_{i=1}^n \mathbb{1}\{h < x_i\} + (+1) \sum_{i=1}^n \mathbb{1}\{h > x_i\} \right]$$

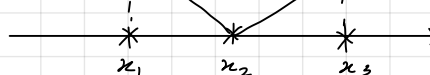
$$\sum_{i=1}^n \mathbb{1}\{h < x_i\} = \# \text{ of } x \text{ greater than } h.$$

$$\sum_{i=1}^n \mathbb{1}\{h > x_i\} = \# \text{ of } x \text{ less than } h.$$

$h^* = \text{Median}(D)$,
best prediction for
mean absolute error.

$$= \frac{1}{n} \left[\underbrace{\sum_{x_i < h} 1}_{\# \text{ of } x \text{ less than } h} - \underbrace{\sum_{x_i > h} 1}_{\# \text{ of } x \text{ greater than } h} \right]$$

odd:
number is unique

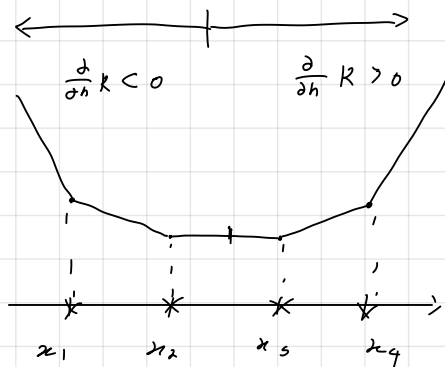


loss function (convex.)

$$\text{loss}(h, x_i) = \mathbb{1}\{h \neq x_i\}$$

$$\text{ERM} \rightarrow \text{Mode}(D)$$

even:
any number between
the middle two data
points minimizes mean
absolute error.



$$< 0 \quad \frac{\partial}{\partial h} K = 0 \quad > 0$$

$$\text{Ex. } R_{\text{abs}}(h, D) = \frac{1}{n} \sum_{i=1}^n |h - x_i|$$

$$R_{\text{sq}}(h, D) = \frac{1}{n} \sum_{i=1}^n (h - x_i)^2$$

$$\sqrt{R_{\text{sq}}(h, D)}$$

$$\text{If } R_{\text{abs}}(h, D) \leq \varepsilon, \quad \Rightarrow \sqrt{R_{\text{sq}}(h, D)} \leq \varepsilon$$

Empirical risk - the average loss on the data sets:

$$R_L(h) = \frac{1}{n} \sum_{i=1}^n L(h, y_i)$$

$\uparrow$
R = "risk"

The goal of learning: find h that minimizes R_L .

This is called empirical risk minimization (ERM).

- The choice of loss function determines the properties of the result.

↓
different loss function
" "
different minimizer
" "
different prediction

ex). $R_{0,1}(h) = \frac{1}{n} \sum_{i=1}^n \begin{cases} 0 & \text{if } h=y \\ 1 & \text{if } h \neq y \end{cases} \Rightarrow h^* = \text{Mode}(P)$

loss	Minimizer	outliers	Differentiable
L_{abs}	median	insensitive	no
L_{sq}	mean	sensitive	yes
$L_{0,1}$	mode	insensitive	no

$$L_{abs}(h, y) = |y - h|$$

$$L_{sq}(h, y) = (y - h)^2$$

Empirical risk $R(h) = \frac{1}{n} \sum_{i=1}^n \text{Loss}(h, y_i)$.

$$h^* = \arg \min_{h \in \mathcal{H}} R(h; D).$$

$$R(h^*; D)$$

$$R_{sq}(h^*; D) = \frac{1}{n} \sum_{i=1}^n (h^* - x_i)^2$$

$$= \frac{1}{n} \sum_{i=1}^n (\bar{x} - x_i)^2 \quad \leftarrow \text{Variance}$$

Standard deviation $\sqrt{R_{sq}(h^*; D)}$

$$\begin{aligned} R_{abs}(h^*; D) &= \frac{1}{n} \sum_{i=1}^n |h^* - x_i| \\ &= \frac{1}{n} \sum_{i=1}^n |\text{Median}(D) - x_i| \end{aligned}$$

$$D = \{1, 2, 2, 4\}$$

$$R_{abs}(h^*; D) = \frac{1}{2}$$

$$\begin{aligned} R_{0,1}(h^*; D) &= \frac{1}{n} \sum_{i=1}^n \mathbb{1}\{h \neq x_i\} \\ &= \frac{n-1}{n} \end{aligned}$$

$$\sqrt{K_{sq}(h; b)} \quad \text{vs.} \quad K_{abs}(h; b)$$

↓

$$K_{sq}(h; b) = \frac{1}{n} \sum_{i=1}^n (h - x_i)^2$$

$$= \frac{1}{n} \sum_{i=1}^n \underbrace{|h - x_i|}_{\parallel}^2$$

$|h - x_i| = a_i$

$$= \frac{1}{n} \sum_{i=1}^n a_i^2$$

$$= \frac{1}{n} \sum_{j=1}^n \left(\frac{1}{n} \sum_{i=1}^n a_i^2 \right)$$

$$= \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n a_i^2$$

$$= \frac{1}{n^2} \sum_{j=1}^n \sum_{i=1}^n a_j^2$$

$$\frac{1}{n^2} \sum_i \sum_j \frac{a_i^2 + a_j^2}{2}$$

$$K_{abs}^2(h; b) = \left(\frac{1}{n} \sum_{i=1}^n |h - x_i| \right)^2$$

$$= \left(\frac{1}{n} \sum_{i=1}^n a_i \right)^2$$

$$= \left(\frac{1}{n} \sum_{i=1}^n a_i \right) \left(\frac{1}{n} \sum_{j=1}^n a_j \right)$$

$$= \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n a_i a_j$$

$$K_{sq}(h; b) - K_{abs}^2(h; b)$$

$$= \frac{1}{n^2} \sum_i \sum_j \left(\frac{a_i^2 + a_j^2}{2} - a_i a_j \right)$$

$$= \frac{1}{n^2} \sum_i \sum_j \left(\frac{1}{2} (a_i^2 + a_j^2 - 2a_i a_j) \right)$$

$$= \frac{1}{n^2} \sum_i \sum_j \left(\frac{1}{2} (a_i - a_j)^2 \right) \geq 0$$

↓

$$K_{sq}(h; b) \geq K_{abs}^2(h; b)$$

↓

$$\sqrt{K_{sq}(h; b)} \geq K_{ab}(h; b), \quad \forall h.$$

Note:

$$\text{Note: } \frac{a^2 + b^2}{2} \geq ab$$

Young's Ineq. S. (a)

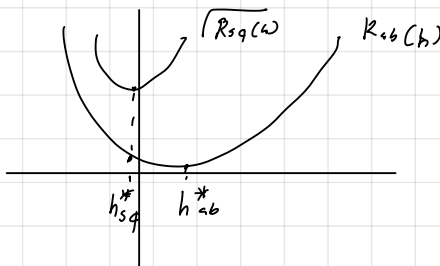
$$\sqrt{K_{sq}(h^*; b)} \quad \text{vs.} \quad K_{ab}(h^*; b)$$

$$\min \sqrt{K_{sq}(h; b)} \geq \min K_{ab}(h; b)$$

$$K_{abs}(h_{ab}^*, b) \leq K_{ab}(h_{sq}^*, b)$$

$$\leq \sqrt{K_{sq}(h_{sq}^*, b)}$$

⇒



Center and Spread

- The input h^* that minimizes $R(h)$ is some measure of the center of the data set.
e.g. median, mean, mode.
- The minimum output $R(h^*)$ represent some measure of the spread, or variance, in the data set.

↓

$$\text{For absolute loss: } R_{abs}(h^*) = R_{abs}(\text{Median}(y_1, y_2, \dots, y_n))$$
$$\uparrow = \frac{1}{n} \sum_{i=1}^n |y_i - \text{median}(y_1, y_2, \dots, y_n)|$$

this is the mean absolute deviation from the median.

$$\text{For squared loss: } R_{sq}(h^*) = R_{sq}(\text{Mean}(y_1, \dots, y_n))$$
$$\uparrow = \frac{1}{n} \sum_{i=1}^n (y_i - \text{Mean}(y_1, \dots, y_n))^2$$

this is called the variance

↑
the square root of variance is standard deviation.

Huber loss / Gradient Descent

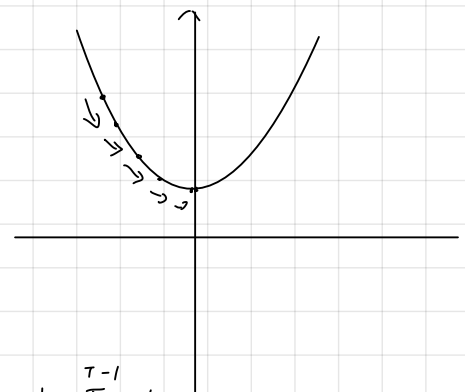
$$\text{loss}(h, x_i) = \begin{cases} |h - x_i| - \frac{1}{4}, & |h - x_i| > \frac{1}{2} \\ (h - x_i)^2, & |h - x_i| \leq \frac{1}{2} \end{cases}$$

$$R_{\text{Huber}} = \frac{1}{n} \sum_{i=1}^n \text{loss}(h, x_i)$$

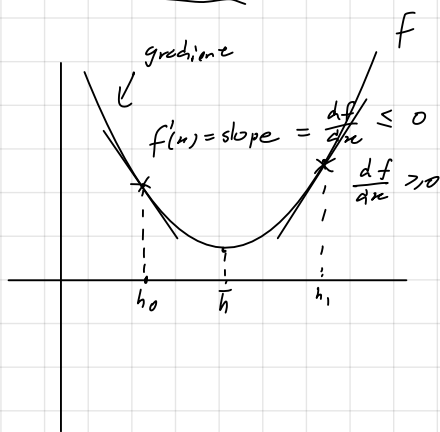
Initialize $h_0 = 0$

$$h_{t+1} = h_t - \Delta t \cdot \frac{d}{dh} R_{\text{Huber}}(h, 0)$$

$$\frac{dh(t)}{dt} = \frac{\partial}{\partial h} R_{\text{Huber}}(h_t, 0) \quad \text{Output } h_t = \frac{1}{t} \sum_{\epsilon=1}^{T-1} h_{\epsilon}$$



Gradient Descent



From Math 101:

gradient $\rightarrow$ direction in which the most

• Pick α to be a positive number, the learning rate.

• Pick a starting prediction, h_0 .

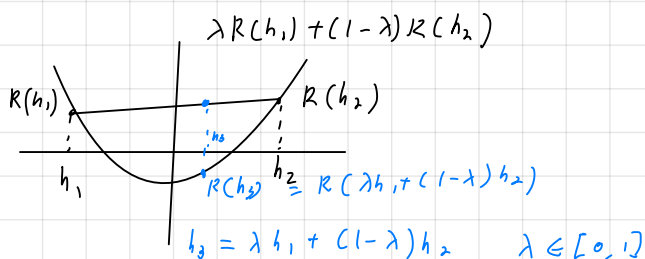
• On step i , perform update $h_i = h_{i-1} - \alpha \cdot \frac{dK}{dh}(h_{i-1})$

• Repeat until convergence.

$$x_{t+1} = x_t - \Delta \frac{df(x_t)}{dx}$$

$$t = 0, 1, 2, \dots$$

Convex function



$$R(h_3) = R(\lambda h_1 + (1-\lambda)h_2)$$

$$R(\lambda h_1 + (1-\lambda)h_2) \leq \lambda R(h_1) + (1-\lambda)R(h_2)$$

$$\forall \lambda \in [0, 1], \forall h_1, h_2.$$

(5) b. loss convex $\Rightarrow R(h, \theta)$ convex.

Thm: For Huber loss, with suitable $\{\Delta_0, \dots, \Delta_T\}$ ^{step size}

$$\bar{h} = \frac{1}{T} \sum_{t=0}^{T-1} h_t$$

all h_t from gradient descent algorithm

$$R(\bar{h}_T, \theta) - R(h^*, \theta) \leq \frac{\epsilon(h_0 - h^*)^2}{T}$$

this means
huber loss will converge

number of iteration that G.D. converges.

If you want error ε :

range of data point

$$\text{Need } T \geq \frac{4(h_0 - h^*)^2}{\varepsilon} = \frac{4(h^*)^2}{\varepsilon} \leq \frac{4K^2}{\varepsilon}$$

$$\Delta t = \frac{1}{4} \quad \forall t. \quad \text{for Huber}$$

$$\text{proof: } \frac{\partial^2}{\partial h^2} \text{loss}(h, h_i) = \begin{cases} 0, & |h - x_i| > \frac{1}{2} \\ 2, & |h - x_i| \leq \frac{1}{2} \end{cases}$$

≥ 0 Huber loss only has minimum, no maximum

$$0 \leq \frac{\partial^2}{\partial h^2} \text{loss}(h, h_i) \leq 2$$

↑
convex
characteristic

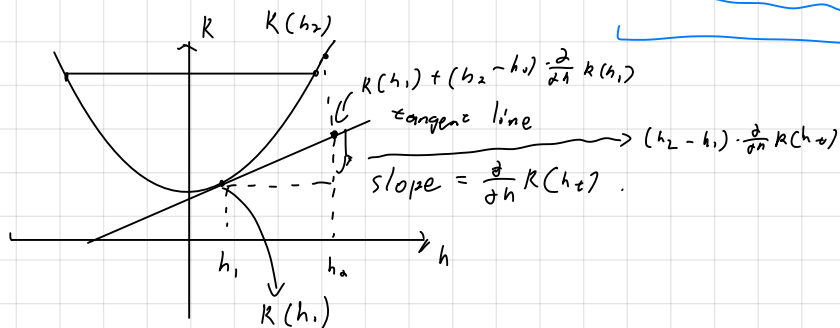
↑
smooth function

$$\Rightarrow 0 \leq \frac{\Delta^2 K(h, b)}{\Delta h^2} \leq 2$$

↖ risk

$$\Rightarrow \left(\frac{\partial}{\partial h} K(h) \right)^2 \leq 4(K(h) - K(h^*)) \quad \text{for } 1D.$$

$$\begin{aligned} (h_{t+1} - h^*)^2 &= \left(h_t - \Delta t \cdot \frac{\partial}{\partial h} K(h_t) - h^* \right)^2 \\ &= \left[(h_t - h^*) - \Delta t \cdot \frac{\partial}{\partial h} K(h_t) \right]^2 \\ &= (h_t - h^*)^2 - 2\Delta t (h_t - h^*) \cdot \frac{\partial}{\partial h} K(h_t) + \Delta t^2 \left(\frac{\partial}{\partial h} K(h_t) \right)^2 \end{aligned}$$



$$\text{So, } K(h_2) \geq K(h_1) + (h_2 - h_1) \frac{\partial}{\partial h} K(h_1)$$

$$\Rightarrow (h_2 - h_1) \cdot \frac{\partial}{\partial h} K(h_1) \geq K(h_1) - K(h_2)$$

$$(h_t - h^*) \cdot \frac{\partial}{\partial h} K(h^*) \geq K(h_t) - K(h^*)$$

$$\leq (h_t - h^*)^2 - 2\Delta t (K(h_t) - K(h^*)) + \Delta t^2 (K(h_t) - K(h^*))$$

1

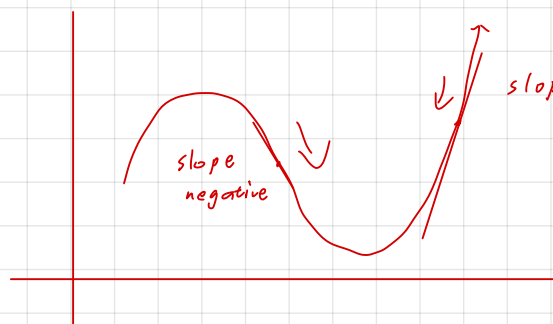
$$\leq (h_t - h^*)^2 - \underbrace{2\Delta t(1-2\Delta t)}_{\text{maximize}} (R(h_t) - R(h^*)).$$

maximize

$\Downarrow$

$$\Delta t = \frac{1}{4}$$

Gradient Descent



$$h_1 = h_0 - \underbrace{\frac{dR}{dh}}_{\text{opposite the direction of the derivative / gradient}}(h_0)$$

Algorithm :

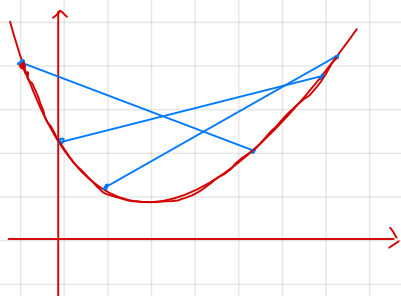
- Pick α to be positive number, the learning rate, also known as step size.
- Pick a starting prediction, h_0 .
- On step i , perform update $h_i = h_{i-1} - \alpha \cdot \frac{dR}{dh}(h_{i-1})$
- Repeat until convergence.

Convex function

• f is convex if for every a, b in the domain of f , the line segment between

$$(a, f(a)) \text{ \& } (b, f(b)).$$

does not go below the plot of f .



function $f: \mathbb{R} \rightarrow \mathbb{R}$ is convex if $\forall a, b \in \mathbb{R}$ and $t \in [0, 1]$,

$$(1-t)f(a) + tf(b) \geq f((1-t)a + tb).$$

Thm: If $R(h)$ is convex and differentiable, then gradient descent converges to a global minimum of R provided that the step size is small enough.

↑

because f convex and has local min $\Rightarrow$ local min is global min.

Test for convexity.

f is convex $\Leftrightarrow \frac{d^2 f}{dx^2}(x) \geq 0 \quad \forall x$.

• Sum of convex functions is convex.

$\hookrightarrow$ if $L(h, y)$ is convex, then $R(h) = \frac{1}{n} \sum_{i=1}^n L(h, y_i)$ is convex.

• $L_{sq}(h, y) = (y - h)^2$ is convex

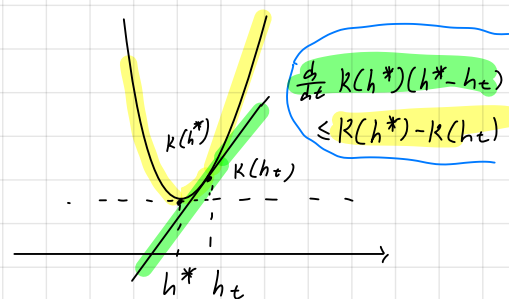
• $L_{abs}(h, y) = |y - h|$ is convex

Gradient Descent

$$h_{t+1} = h_t - \alpha \frac{d}{dh} R(h_t)$$

$$0 \leq \frac{d^2}{dh^2} R(h) \leq L \leftarrow \begin{matrix} \uparrow & \uparrow \\ \text{convex} & \text{smooth} \end{matrix}$$

$L=2$ for Huber Risk



$$\begin{aligned} (h_{t+1} - h^*)^2 &= (h_t - \alpha \frac{d}{dh} R(h_t) - h^*)^2 \\ &= (h_t - h^*)^2 - 2 \alpha \underbrace{\frac{d}{dh} R(h_t) (h_t - h^*)}_{\geq R(h^*) - R(h_t)} + \alpha^2 \underbrace{\left(\frac{d}{dh} R(h_t) \right)^2}_{\text{HW 2 QS}} \\ &\leq (h_t - h^*)^2 - 2 \alpha (1 - L \cdot \alpha) (R(h_t) - R(h^*)). \end{aligned}$$

$\alpha = \frac{1}{2L}$ for optimize
for maximum $\Rightarrow$ less iteration for gradient descent.

So, let $\alpha = \frac{1}{2L}$,

$$\Rightarrow (h_{t+1} - h^*)^2 \leq (h_t - h^*)^2 - \frac{1}{2L} (R(h_t) - R(h^*))$$

this gets 0 when $R(h_t) = R(h^*)$, in which case, converges.

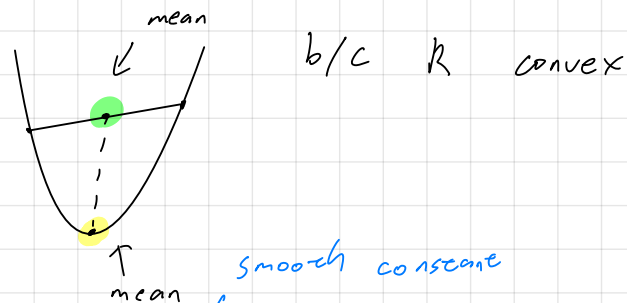
$$\Rightarrow \frac{1}{2L} (R(h_t) - R(h^*)) \leq (h_t - h^*)^2 - (h_{t+1} - h^*)^2$$

$$\text{Sum} \left\{ \begin{aligned} \frac{1}{2L} (R(h_0) - R(h^*)) &\leq (h_0 - h^*)^2 - (h_1 - h^*)^2 \\ \frac{1}{2L} (R(h_1) - R(h^*)) &\leq (h_1 - h^*)^2 - (h_2 - h^*)^2 \\ &\vdots \\ \frac{1}{2L} (R(h_{t-1}) - R(h^*)) &\leq (h_{t-1} - h^*)^2 - (h_t - h^*)^2 \end{aligned} \right. \quad \begin{matrix} \text{telescoping sum} \\ \text{(canceling out terms)} \end{matrix}$$

$$\Rightarrow \frac{1}{2L} \sum_{t=0}^{T-1} (R(h_t) - R(h^*)) \leq (h_0 - h^*)^2 - (h_T - h^*)^2 \leq (h_0 - h^*)^2$$

$$\Rightarrow \frac{1}{2L} \cdot \frac{1}{T} \sum_{t=0}^{T-1} (R(h_t) - R(h^*)) \leq \frac{1}{T} (h_0 - h^*)^2$$

let $\bar{h}_t = \frac{1}{T} \sum_{t=0}^{T-1} h_t$, then $K(\bar{h}_t) - K(h^*) \leq \frac{1}{T} \sum_{t=0}^{T-1} K(h_t) - K(h^*)$.



$$K(\bar{h}_t) - K(h^*) \leq \frac{2L}{T} (h_0 - h^*)^2$$

If we want

$$K(\bar{h}) - K(h^*) \leq \varepsilon, \quad \forall \varepsilon > 0.$$

$$\text{take } T = \frac{4}{\varepsilon} (h_0 - h^*)^2.$$

Bound $|\bar{x}(D_5^{bs}) - \bar{x}(D_\infty^{vs})|$

$$(\bar{x}(D_n) - \bar{x}(D_\infty))^2 = O(1/n)$$

how large constant

$$(\bar{x}(D_1) - \bar{x}(D_\infty))^2$$

Use D_n to estimate: $\leq \frac{C}{n} \sum_{i=1}^n (x_i - \bar{x}(D_\infty))^2$ $\leq \frac{C}{n} \sum_{i=1}^n (x_i - \bar{x}(D_n))^2$

true mean empirical mean

$$\sum_{i=1}^n (x_i - \bar{x}(D_\infty))^2 \approx \sum_{i=1}^n (x_i - \bar{x}(D_n))^2$$

||

$$= C \cdot \text{Var}(D_n)$$

$$\begin{aligned} & (x_i - \bar{x}(D_\infty) + \bar{x}(D_\infty) - \bar{x}(D_n))^2 \\ &= (x_i - \bar{x}(D_\infty))^2 + (\bar{x}(D_\infty) - \bar{x}(D_n))^2 + 2(\dots)(\dots) \\ & \quad \underbrace{\hspace{10em}} \\ & \quad \frac{1}{n} \cdot (x_i - \bar{x}(D_\infty))^2 \end{aligned}$$

Put together: $(\bar{x}(D_n) - \bar{x}(D_\infty))^2 \leq \frac{C}{n} \cdot \text{Var}(D_n)$

$$|\bar{x}(D_n) - \bar{x}(D_\infty)| \leq \frac{C}{\sqrt{n}} \cdot \text{STD}(D_n)$$

$$|\bar{x}(D_n) - \bar{x}(D_0)| \leq \frac{\log(1/\delta)}{\sqrt{n}} \cdot \text{STD}(D_n)$$

Sequence Decision Making (Reinforcement Learning)

- Bandits Problem.

↓

Explore then commit

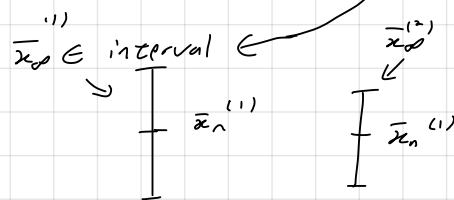
- spend the first N rounds uniformly on the two options.
- Then commit to the one with higher average reward: $\arg\min_{a=1,2} \bar{R}^a$

$$\text{Want: } \Delta \text{reward} = O\left(\frac{\text{STD}(D_N)}{\sqrt{n}}\right)$$

↑
confidence level constant

$$\Leftrightarrow |\bar{x}_n - \bar{x}_0| \leq \frac{\text{STD}(D_n)}{\sqrt{n}} \cdot \log\left(\frac{1}{\delta}\right)$$

↑
(1-δ) probability



How to compare $\Rightarrow$ when two interval do not overlap

⇓

one is definitely greater than the other.

$$\Uparrow \text{ by let } \log\left(\frac{1}{\delta}\right) = |\bar{x}_0^{(1)} - \bar{x}_0^{(2)}|$$

= Δreward

↑
difference in true mean

How to figure out?

Algorithm:

- keep track N^a

$$\text{• optimistic estimate: } \hat{R}^a = \bar{R}^a + C \cdot \frac{\text{STD}(D_{N^a})}{N^a}$$

- choose action $\arg\min_{a=1,2} \hat{R}^a$

↑
of times
chosen action a

Linear Regression

- Feature - an attribute, information
 - predictor variable - the feature - x
 - response variable - the quantity - y - we are trying to predict.
 - hypothesis function/prediction rule.
- Use loss function to quantify the quality of prediction.

- Squared loss : $(\text{actual} - \text{predicted})^2$

↓

Empirical risk : $R_{sq}(H) = \frac{1}{n} \sum_{i=1}^n (y_i - H(x_i))^2$

↓

Smallest mean squared error $\rightarrow$ good prediction

↑

Not overfitting!

↓

Choose H to be certain function!

↓

linear function : $H(x) = w_0 + w_1 x$

↑ ↑
intercept slope

↓

H^* be the linear function that minimizes $R_{sq}(H) = \frac{1}{n} \sum_{i=1}^n (y_i - H(x_i))^2$ parameter
(least squares regression).

↓ plug in H

$$R_{sq}(w_0, w_1) = \frac{1}{n} \sum_{i=1}^n (y_i - (w_0 + w_1 x_i))^2$$

Solve for gradient

↑

Find slope and intercept that minimizes.

↓

$$\frac{\partial R_{sq}}{\partial w_0} = -\frac{2}{n} \sum_{i=1}^n (y_i - (w_0 + w_1 x_i))$$

$$\frac{\partial R_{sq}}{\partial w_1} = -\frac{2}{n} \sum_{i=1}^n (y_i - (w_0 + w_1 x_i)) x_i$$

↓

Solve for $\frac{\partial R_{sq}}{\partial w_0} = 0$ & $\frac{\partial R_{sq}}{\partial w_1} = 0$

$$w_0 = \frac{1}{n} \sum_{i=1}^n y_i - w_1 \frac{1}{n} \sum_{i=1}^n x_i$$

$$w_0^* = \bar{y} - w_1^* \bar{x}$$

↑
best intercept

$$w_1^* = \frac{\sum_{i=1}^n (y_i - \bar{y}) x_i}{\sum_{i=1}^n (x_i - \bar{x}) x_i}$$

where

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$$

$$\bar{y} = \frac{1}{n} \sum_{i=1}^n y_i$$

↑
best slope

$$= \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

the process is called
"fitting to the data".

least square solution
(optimal parameter).

$$H^*(w) = w_0^* + w_1^* x$$

• Correlation coefficient

- a measure of the strength of the linear association of two variables, x & y .
- coefficient, r , the average of the product of x and y , in standard unit

$$r = \frac{1}{n} \sum_{i=1}^n \left(\frac{x_i - \bar{x}}{s_x} \right) \left(\frac{y_i - \bar{y}}{s_y} \right)$$

→ mean
→ standard deviation.

↓

- r has no unit
- $r = 1$ → perfect positive linear association
- $r = -1$ → perfect negative linear association
- the closer r is to 0 → the weaker the linear association

↓ optimal slope for the linear prediction rule w_1^* in r

$$w_1^* = r \frac{s_y}{s_x} \rightarrow \bullet \text{ sign}(r) = \text{sign}(w_1^*)$$

• intercept

$$H^*(\bar{x}) = \bar{y}$$

• y value get spread out → $s_y \uparrow$ → slope $\uparrow$

• x value get spread out → $s_x \uparrow$ → slope $\downarrow$

Regression

ex. Look at salary (y) as function of years of experience (x).

↑
response variable

↑
predictor variable

$$D = \{ (x_1, y_1), \dots, (x_n, y_n) \}$$

$$y \approx H(x)$$

↑
Hypothesis function / prediction rule.

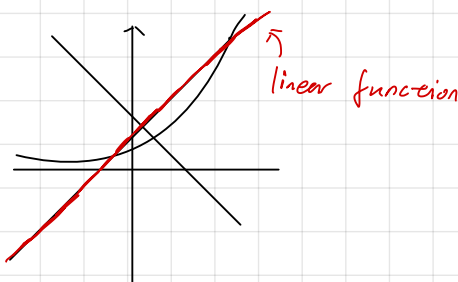
Ex. $y = ax + b$

$$y = e^{ax}$$

$$y = b - x$$

... ..

which $H(x)$ is good?



loss function

$$\text{loss}(H; (x_i, y_i)) = (H(x_i) - y_i)^2 \quad \text{squared loss}$$

$$\begin{aligned} \text{Empirical risk: } K_{sq}(H, D) &= \frac{1}{n} \sum_{i=1}^n \text{loss}(H; (x_i, y_i)) \\ &= \frac{1}{n} \sum_{i=1}^n (H(x_i) - y_i)^2 \end{aligned}$$

Optimization: choose prediction rule by $\min_{H \in \mathcal{H}} K_{sq}(H; D)$

One of all loss function: $H(x) = w_1 x + w_0$. Find H^* that minimizes $K_{sq}(H; D)$.

squared loss ↷

↓
optimize w_1 & w_0 .

$$\begin{aligned} \min_{w_1, w_0 \in \mathbb{R}} K_{sq}(H; D) \\ &= \min \frac{1}{n} \sum_{i=1}^n (H(x_i) - y_i)^2 \\ &= \min \frac{1}{n} \sum_{i=1}^n (w_1 x_i + w_0 - y_i)^2 \end{aligned}$$

optimal prediction

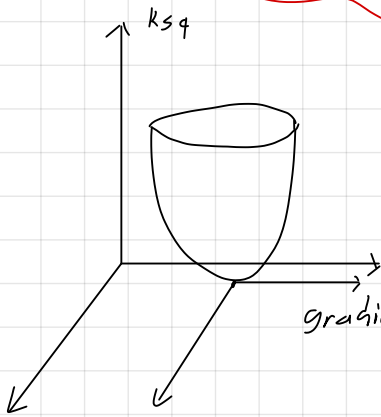
$$w_1^*, w_0^* = \underset{w_1, w_0}{\operatorname{argmin}} \frac{1}{n} \sum_{i=1}^n (w_1 x_i + w_0 - y_i)^2$$

$$\text{let } \nabla_{(w_1, w_0)} \left(\frac{1}{n} \sum_{i=1}^n (w_1 x_i + w_0 - y_i)^2 \right) = 0$$

⇓

$$\nabla_{(w_1, w_0)} (f(w_1, w_0)) = \begin{pmatrix} \frac{\partial f}{\partial w_1} \\ \frac{\partial f}{\partial w_0} \end{pmatrix} = 0$$

solve for critical point



gradient / partial derivative = 0.

$$\Downarrow$$

$$\begin{cases} \frac{\partial Ksq(w_1, w_0, D)}{\partial w_0} = 0 & ① \\ \frac{\partial Ksq(w_1, w_0, D)}{\partial w_1} = 0 & ② \end{cases}$$

$$① \frac{\partial}{\partial w_0} \left(\frac{1}{n} \sum_{i=1}^n (w_1 x_i + w_0 - y_i)^2 \right) = 0$$

$$= \frac{1}{n} \sum_{i=1}^n \frac{\partial}{\partial w_0} [(w_1 x_i + w_0 - y_i)^2]$$

$$= \frac{1}{n} \sum_{i=1}^n 2 \cdot (w_1 x_i + w_0 - y_i)$$

$$= 2 \left(\frac{1}{n} \sum_{i=1}^n w_0 + \frac{1}{n} \sum_{i=1}^n w_1 x_i - \frac{1}{n} \sum_{i=1}^n y_i \right)$$

$$= 2 (w_0 + w_1 \bar{x} - \bar{y})$$

$\bar{x} = x \text{ mean}$
 $\bar{y} = y \text{ mean}$

$\Downarrow$

$$2(w_0 + w_1 \bar{x} - \bar{y}) = 0$$

$w_1 \bar{x} + w_0 = \bar{y}$ optimal solution pass through $(\bar{x}, \bar{y})$.

$$w_0 = \bar{y} - w_1 \bar{x}$$

$$w_1 \frac{1}{n} \sum_{i=1}^n x_i (x_i - \bar{x})$$

$\Downarrow$

$$w_1 \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(x_i - \bar{x})$$

$$= w_1 \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2$$

$$\text{since } \frac{1}{n} \sum_{i=1}^n \bar{x} (x_i - \bar{x}) = 0$$

$\Downarrow$

$$② \frac{\partial}{\partial w_1} \left(\frac{1}{n} \sum_{i=1}^n (w_1 x_i + w_0 - y_i)^2 \right) = 0$$

$$= \frac{1}{n} \sum_{i=1}^n \left(\frac{\partial}{\partial w_1} (w_1 x_i + w_0 - y_i)^2 \right)$$

$$= \frac{1}{n} \sum_{i=1}^n (2 \cdot (w_1 x_i + w_0 - y_i) \cdot x_i)$$

$$= 2 \frac{1}{n} \sum_{i=1}^n (w_1 x_i^2 + w_0 x_i - x_i y_i)$$

② + ①:

$$0 = \frac{1}{n} \sum_{i=1}^n (w_1 x_i^2 + x_i (\bar{y} - w_1 \bar{x}) - x_i y_i)$$

$$0 = \frac{1}{n} \sum_{i=1}^n (w_1 x_i^2 + x_i \bar{y} - x_i w_1 \bar{x} - x_i y_i)$$

$$0 = w_1 \frac{1}{n} \sum_{i=1}^n (x_i^2 - x_i \bar{x}) + \frac{1}{n} \sum_{i=1}^n (x_i \bar{y} - x_i y_i)$$

$$w_1 \frac{1}{n} \sum_{i=1}^n (x_i^2 - x_i \bar{x}) = \frac{1}{n} \sum_{i=1}^n (x_i y_i - x_i \bar{y})$$

$\Downarrow$

$$\frac{1}{n} \sum_{i=1}^n x_i (y_i - \bar{y})$$

$\Downarrow$

$$\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})$$

since

$$\frac{1}{n} \sum_{i=1}^n \bar{x} (y_i - \bar{y}) = 0$$

$$\text{Since } \frac{1}{n} \sum_{i=1}^n (y_i - \bar{y}) = 0$$

sum of products
↓

$$\text{Thus, } w_1 = \frac{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2}$$

↑
variances in x

correlation

Determine
the direction of the
regression line.
Positive or Negative

since

$$\sum \bar{x}(y_i - \bar{y}) = 0$$

$$\sum \bar{y}(x_i - \bar{x}) = 0$$

$$\sum (x_i - \bar{x})(y_i - \bar{y})$$

$$= \sum x_i(y_i - \bar{y}) = \sum x_i y_i - \sum x_i \bar{y}$$

$$= \sum y_i(x_i - \bar{x}) = \sum y_i x_i - \sum y_i \bar{x}$$

⇓

$$= \sum x_i y_i - \bar{x} \bar{y}$$

Midterm Review

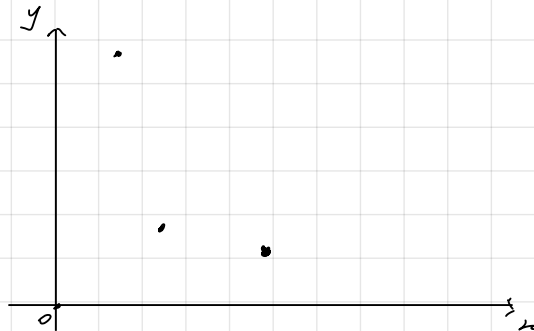
HW1

↑
some question in

$$\bar{x} = 5$$

$$\bar{y} = 4$$

ex)	x_i	y_i
(x_1, y_1)	3	7
(x_2, y_2)	4	3
(x_3, y_3)	8	2



$$y = H(x) = w_1 x + w_0$$

What is the best prediction rule?

that describe the data set?

$$\begin{aligned} w_1^* &= \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2} \\ &= \frac{(-2)(3) + (-1)(-1) + (3)(-2)}{(-2)^2 + (-1)^2 + (3)^2} \\ &= \frac{-6 + 1 - 6}{4 + 1 + 9} \\ &= \frac{-11}{14} \end{aligned}$$

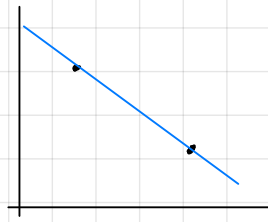
$$\begin{aligned} w_0^* &= \bar{y} - w_1^* \bar{x} \\ &= 4 - \frac{-11}{14} \cdot 5 \\ &= \frac{101}{14} \end{aligned}$$

$$\text{So, } H(x) = -\frac{11}{14}x + \frac{101}{14}$$

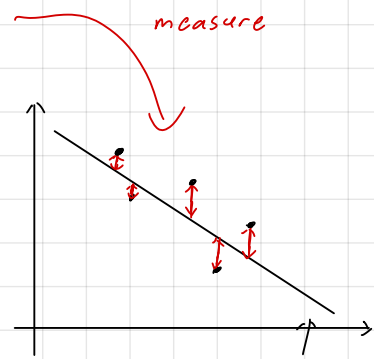
$Rsq(H^*; D) = Rsq(w_1^*x + w_0^*; D)$. Mean squared error of a data set

$$= \frac{1}{n} \sum_{i=1}^n ((w_1 * x_i + w_0) - y_i)^2$$

D has two data points

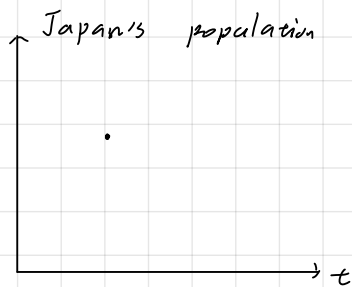


$$R_{sq} = 0.$$



linear prediction rule.

ex). Japan's population decrease by half by 2060, will vanish by 2100.



Choices of function close of H.

1. Linear
2. Polynomial
3. Exponential.

next year population

Current population

$$X_{t+1} = X_t + \alpha \cdot X_t$$

birth rate

$$X_{t+1} = X_t + \alpha \cdot X_t - \beta X_t$$

$$= (1 + \alpha - \beta) X_t$$

death rate

$$= \delta X_t$$

$$X_t = \delta X_{t-1} = \delta^2 \cdot X_{t-2} = \dots = \delta^t \cdot X_0$$

initial population

$$X(t) = \delta^t \cdot w$$

Can we learn (δ, w) from data $\{(t_1, x_1) \dots (t_n, x_n)\}$?

$$\ln(X(t)) = \ln(\delta^t \cdot w)$$

$$= \ln(\delta^t) + \ln(w)$$

$$= t \cdot \ln(\delta) + \ln(w)$$

w_1

w_0

$$\Downarrow$$

$$\ln(x(t)) = w_1 t + w_0$$

$$\begin{matrix} \text{So, } & t_i & \ln(x(t_i)) \\ & \vdots & \vdots \\ & \vdots & \vdots \end{matrix}$$

do least square solution
for w_1 and w_0 .

Conclusion:

$$g(y) = w_0 + w_1 f(x)$$

$$y = f^T \cdot w$$

idea: Transform non-linear relationship to linear one. Then do least square solution.

Multi-linear regression

$$y = w_0 + w_1 x + w_2 x^2 + w_3 x^3$$

Want to fit:

$$y = H(x) = w_0 + w_1 x + w_2 x^2 + w_3 x^3 + \dots + w_n x^n$$

$$y = H(x^{(1)}, x^{(2)}, x^{(3)}, \dots, x^{(n)}) \leftarrow \text{multiple feature}$$

$\uparrow$ $\uparrow$ $\uparrow$ $\uparrow$
 Salary YOE location job family, ...

$$= w_0 + w_1 x^{(1)} + \dots + w_n x^{(n)}$$

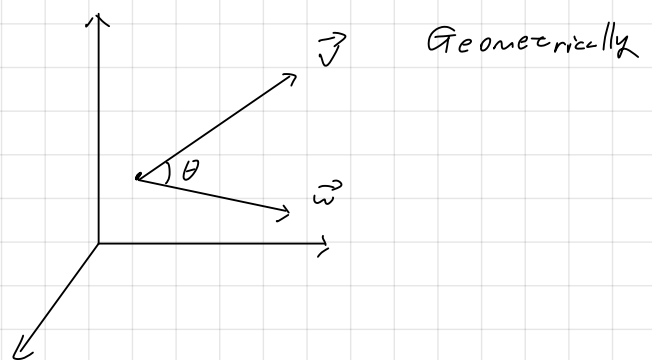
$$= w_0 + \sum_{i=1}^n w_i \cdot x^{(i)}$$

Linear Algebra (Review):

Vector:

$$\vec{v} = \begin{bmatrix} v_1 \\ \vdots \\ v_d \end{bmatrix}$$

$$\vec{v} \in \mathbb{R}^d$$



$$\|\vec{v}\|, \quad \frac{v^T w}{\|v\| \|w\|} = \cos \theta.$$

$$v^T w = \sum_{i=1}^d v_i \cdot w_i = w^T v \leftarrow \text{inner product.}$$

$$[v_1, \dots, v_d] \cdot \begin{bmatrix} w_1 \\ \vdots \\ w_d \end{bmatrix}$$

Matrix :

$$A = \begin{pmatrix} \vec{u}_1^T \\ \vdots \\ \vec{u}_n^T \end{pmatrix} = \begin{pmatrix} u_{11} & \dots & u_{1m} \\ \vdots & \ddots & \vdots \\ u_{n1} & \dots & u_{nm} \end{pmatrix} \quad A \in \mathbb{R}^{n \times m}$$

$$= (\vec{v}_1, \dots, \vec{v}_m).$$

Matrix Vector:

$$\vec{y} = A \cdot w = \begin{bmatrix} u_{11} & \dots & u_{1m} \\ \vdots & \ddots & \vdots \\ u_{n1} & \dots & u_{nm} \end{bmatrix} \begin{bmatrix} w_1 \\ \vdots \\ w_m \end{bmatrix}$$

$\begin{matrix} \mathbb{R}^{n \times m} & \downarrow & \mathbb{R}^m \\ & \mathbb{R}^n & \end{matrix}$

$$= \begin{bmatrix} u_1^T \cdot w \\ \vdots \\ u_n^T \cdot w \end{bmatrix} \in \mathbb{R}^n$$

$$= \sum_{i=1}^m w_i \cdot \vec{v}_i$$

$$\vec{y} \in \text{span} \{ \vec{v}_1, \dots, \vec{v}_m \}$$

Using L.A. for Multiple linear regression:

$$s_0, \quad w_0 + \sum_{i=1}^d w_i \cdot x^{(i)} \quad \downarrow \quad \vec{w}^T \cdot \vec{x}$$

$$\vec{w} = \begin{bmatrix} w_1 \\ \vdots \\ w_d \end{bmatrix}, \quad \vec{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_d \end{bmatrix}$$

$$= w_0 x^{(0)} + \sum_{i=1}^d w_i \cdot x^{(i)}, \quad \text{where } x^{(0)} = 1.$$

$$= \sum_{i=0}^d w_i \cdot x^{(i)}$$

$$= \vec{x}^T \cdot \vec{w}$$

$$\Leftarrow \text{Thus, } \vec{w} = \begin{bmatrix} w_0 \\ w_1 \\ \vdots \\ w_d \end{bmatrix}, \quad \vec{x} = \begin{bmatrix} 1 \\ x_1 \\ \vdots \\ x_d \end{bmatrix}$$

$$H(\vec{x}) = \vec{x}^T \vec{w}$$

$$\vec{x} = \begin{pmatrix} 1 \\ x_1 \\ \vdots \\ x_d \end{pmatrix} \in \mathbb{R}^{d+1}$$

$$\vec{w} = \begin{pmatrix} w_0 \\ \vdots \\ w_d \end{pmatrix} \in \mathbb{R}^{d+1}$$

$$\vec{w}, \vec{x} \in \mathbb{R}^{(d+1)}$$

Then, $y = H(\vec{x}) = \vec{x}^T \cdot \vec{w}$ | A data set will be like

parameter to learn	Salary	1	YOE	location	...
	y_1	1	$x_1^{(1)}$	$x_1^{(2)}$	...
	y_2	1	$x_2^{(1)}$	$x_2^{(2)}$	
	$\vdots$	$\vdots$	$\vdots$	$\vdots$	
	y_n	1	$x_n^{(1)}$	$x_n^{(2)}$	

Data point:
 $(\vec{x}, y)$.

Least Square:

$$\text{Loss}(H(\vec{x}, y)) = (y_i - H(\vec{x}_i))^2$$

$$\text{Loss}(\vec{w}, (\vec{x}_i, y)) = (y_i - \vec{x}_i^T \vec{w})^2$$

$$\text{So, } R(\vec{w}, D) = \frac{1}{n} \sum_{i=1}^n \text{Loss}(\vec{w}, (\vec{x}_i, y_i))$$

$$D = \{(\vec{x}_1, y_1), \dots, (\vec{x}_n, y_n)\}$$

$$= \frac{1}{n} \sum_{i=1}^n (y_i - \vec{x}_i^T \vec{w})^2$$

$$= \frac{1}{n} \|y - \vec{X}^T \vec{w}\|^2 \quad \leftarrow \| \vec{v} \|^2 = \vec{v}^T \vec{v} = \sum_{i=1}^n v_i^2$$

$$= \frac{1}{n} \| \vec{V} \|^2 \quad \text{let } \underline{V_i = y_i - \vec{x}_i^T \vec{w}}$$

$$\text{So, } \vec{V} = \begin{bmatrix} y_1 - \vec{x}_1^T \vec{w} \\ y_2 - \vec{x}_2^T \vec{w} \\ \vdots \\ y_n - \vec{x}_n^T \vec{w} \end{bmatrix} = \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} - \begin{bmatrix} \vec{x}_1^T \vec{w} \\ \vdots \\ \vec{x}_n^T \vec{w} \end{bmatrix} = \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} - \vec{w} \begin{bmatrix} \vec{x}_1 \\ \vdots \\ \vec{x}_n \end{bmatrix}$$

$$= \vec{y} - \vec{X} \cdot \vec{w}$$

$\begin{matrix} \uparrow & & \uparrow \\ \mathbb{R}^n & & \mathbb{R}^{d+1} \\ & \uparrow & \\ & \mathbb{R}^{n \times (d+1)} & \end{matrix}$

$\vec{X} \in \mathbb{R}^{n \times (d+1)}$
Design Matrix

Thus, $y = H(\vec{w}) = \vec{w}^T \cdot \vec{w}$ X $\in \mathbb{R}^{n \times (d+1)}$

Empirical Risk

$R(\vec{w}; D) = \frac{1}{n} \|\vec{y} - X \cdot \vec{w}\|^2$

$\vec{y} \in \mathbb{R}^n$
 $\vec{w} \in \mathbb{R}^{d+1}$

$\vec{y} = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}$ $X = \begin{pmatrix} x_1^T \\ \vdots \\ x_n^T \end{pmatrix}$

$= \frac{1}{n} (\vec{y} - X \cdot \vec{w})^T (\vec{y} - X \cdot \vec{w})$

$= \frac{1}{n} (\vec{y}^T - \vec{w}^T X^T) (\vec{y} - X \cdot \vec{w})$

$= \frac{1}{n} (\vec{y}^T \vec{y} - \underbrace{\vec{y}^T X \vec{w}}_{\text{scalar}} - \underbrace{\vec{w}^T X^T \vec{y}}_{\text{scalar}} + \vec{w}^T X^T X \vec{w})$

$= \begin{pmatrix} x_{1(1)} & \dots & x_{1(d)} \\ \vdots & & \vdots \\ x_{n(1)} & \dots & x_{n(d)} \end{pmatrix}$

↑
Design
Matrix

$\vec{w} = \arg \min_{\vec{w} \in \mathbb{R}^{d+1}} R(\vec{w}, D)$

$\vec{y}^T X \vec{w} = (\vec{y}^T X \vec{w})^T = \vec{w}^T X^T \vec{y}$

$= \frac{1}{n} (\vec{y}^T \vec{y} - 2 \vec{w}^T X^T \vec{y} + \vec{w}^T X^T X \vec{w})$

So, $\nabla_{\vec{w}} R_{sq}(\vec{w}, D)$

$\downarrow$
 $\mathbb{R}^1$

$\downarrow$
 $\mathbb{R}^1$

$\downarrow$
 $\mathbb{R}^1$

$= \frac{1}{n} (\underbrace{\vec{y}^T \vec{y}}_{=0} - 2 \underbrace{\vec{w}^T X^T \vec{y}}_{= X^T \vec{y} \in \mathbb{R}^{d+1}} + \underbrace{\vec{w}^T X^T X \vec{w}}_{= 2 X^T X \vec{w}})$

$=$

Since $\nabla_{\vec{w}} \vec{w}^T \vec{b}$

$= \begin{pmatrix} \frac{\partial}{\partial w_0} \\ \vdots \\ \frac{\partial}{\partial w_d} \end{pmatrix} (w_0 b_0 + \dots + w_d b_d)$

$= \begin{pmatrix} b_0 \\ \vdots \\ b_d \end{pmatrix} = \vec{b}$

from HW 4

Thus, $0 = \frac{1}{n} (-2 X^T \vec{y} + 2 X^T X \vec{w})$

$0 = -X^T \vec{y} + X^T X \vec{w}$

$X^T \vec{y} = X^T X \vec{w}$

Normal Equation

$\begin{matrix} & X^T X & \\ \uparrow & & \uparrow \\ \mathbb{R}^{(d+1) \times n} & & \mathbb{R}^{n \times (d+1)} \\ \downarrow & & \\ \mathbb{R}^{(d+1) \times (d+1)} & & \end{matrix}$

invertible

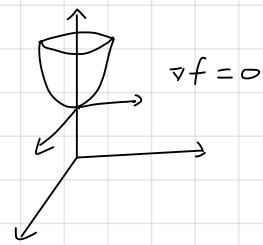
square matrix & full rank

(column of $X^T X$ linearly independent)

From one-dimensional LS

$\begin{cases} \frac{\partial}{\partial w_0} R(w_0, w_1, D) = 0 \\ \frac{\partial}{\partial w_1} R(w_0, w_1, D) = 0 \end{cases}$

multi-dimensional w.r.t vector



$\begin{pmatrix} \frac{\partial}{\partial w_0} R(\vec{w}, D) \\ \vdots \\ \frac{\partial}{\partial w_d} R(\vec{w}, D) \end{pmatrix} = 0$

$\begin{pmatrix} \frac{\partial}{\partial w_0} \\ \vdots \\ \frac{\partial}{\partial w_d} \end{pmatrix} R(\vec{w}, D) = 0$

$\nabla_{\vec{w}} R(\vec{w}, D) = 0$

$\begin{matrix} \uparrow & \uparrow \\ \in \mathbb{R}^{d+1} & \mathbb{R} \end{matrix}$

Thus, $\exists (X^T X)^{-1}$,
 $(X^T X)^{-1} X^T X = I$

Our normal equation becomes

$$\vec{w} = (X^T X)^{-1} X^T \vec{y}$$

why matter?
 $(X^T X) \vec{w} = \sum_{i=0}^d w_i \cdot \vec{b}_i$

if dependent, $\exists \vec{w}' \neq 0$, s.t. $(X^T X) \vec{w}' = 0$.

linear combination of $X^T X$
 s.t. $= 0$.

Q: When is $X^T X$ full rank?

① X is square matrix, $n = d+1$

$$X^T X = \begin{pmatrix} 1 & \dots & 1 \\ x_{1(1)} & \dots & x_{n(1)} \\ \vdots & & \vdots \\ x_{1(d)} & \dots & x_{n(d)} \end{pmatrix} \begin{pmatrix} 1 & x_{1(1)} & \dots & x_{n(1)} \\ \vdots & \vdots & & \vdots \\ 1 & x_{1(d)} & \dots & x_{n(d)} \end{pmatrix}$$

$X^T (X \vec{w}) \rightarrow X^T$ full rank $\rightarrow X$ full rank

(data point not repeating
 or not linearly combination of other data point)

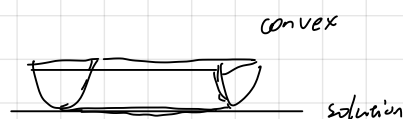
② data point linearly dependent

$\Downarrow$

$n < d+1$ over-parameterize

$\Downarrow$

infinity many solutions.



always solution to
 normal equation

Meaning of optimal solution

$$\vec{y} \approx H(\vec{x}) = \vec{X}^T \vec{w}^*$$

← average salary of industry

$$= w_0^* + w_1^* x_{(1)} + w_2^* x_{(2)} + w_3^* x_{(3)} + \dots + w_d^* x_{(d)}$$

↑
Salary

↑
Y.O.E.

↑
years of between job

Standardize (z-score)

$$\frac{X_{(1)} - \bar{X}_{(1)}}{STD(DC1)}$$

↓

$$STD = \sqrt{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2}$$

If we standardize ↓

w^* reflect correlation.

$$w_3^* : 100,000 \rightarrow 100,001$$

$$w_1^* : 1 \text{ year} \rightarrow 2 \text{ years}$$

↓

$$w_1^* \gg w_3^*$$

$$w_3^* : \$100 \rightarrow \$101$$

scalar

ex) $H(x) = w_0 + w_1 x + w_2 x^2 + w_3 x^3 + \dots + w_d x^d$ ← linear regression okay

limitation of linear regression?

No, b/c Taylor series idea → polynomial approximate any function.

$$H(x) = w_0 + w_1 \sin(x) + d_1 \cos(x) + w_2 \sin(2x) + d_2 \cos(2x) + \dots$$

↑

can be used to model periodical function.

Problem: • Overfitting
• efficiency

ex). Free falling obj.

$$H(x) = w_0 + w_1 x + w_2 x^2 + \dots$$

not necessary

$$H(\vec{x}) = w_0 + w_1^T \vec{x} + \underbrace{\vec{x}^T W_2}_{\text{hard to solve (high dimensional)}} \vec{x} \dots$$

Hierarchical model (deep learning model).

classification problem

0

1

$$h^1(\vec{x}) = \sigma(W^1 \vec{x})$$

$W^1 \in \mathbb{R}^{d \times d}$
depends on how large the next layer you want

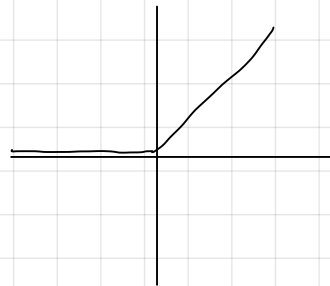
$$h^2(\vec{x}) = \sigma(W^2 h^1(\vec{x}))$$

$$\vdots$$

$$h^t(\vec{x}) = \sigma(W^t h^{t-1}(\vec{x}))$$

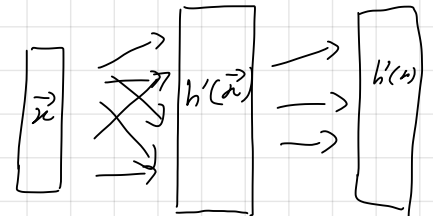
$$\vdots$$

$$h^T(\vec{x}) = \underline{\underline{W^T h^{T-1}(\vec{x})}}$$



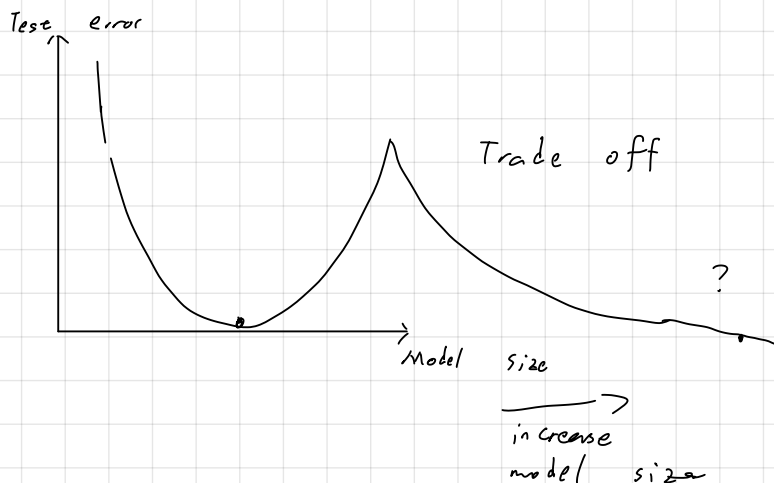
$$\sigma_i(\vec{z}) = \begin{cases} 0 & , z_i < 0 \\ z_i & , z_i \geq 0 \end{cases}$$

Rectified Linear Unit

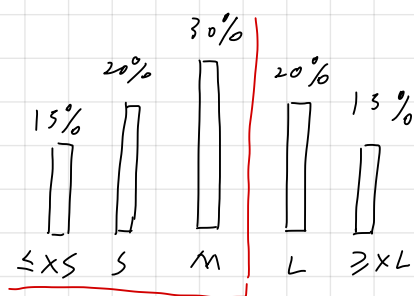


On Final

Vector Calculus



Probability theory



distribution are similar

Probability Statement:

Outcome ("≤XS", "S", "M", "L", "≥XL")

Sample space $S = \{ \downarrow \}$

Event E is a subset of sample space

Probability $\mathbb{P}$: a function that maps event to a real number.

$$\textcircled{1} \quad 0 \leq \mathbb{P}(E) \leq 1, \quad \forall E \subseteq S$$

$$\textcircled{2} \quad \mathbb{P}(E_1 \cup E_2) = \mathbb{P}(E_1) + \mathbb{P}(E_2), \quad \forall E_1 \text{ and } E_2 \text{ not intersecting each other.}$$

$$\forall E_1, E_2 = \emptyset$$

$$\begin{aligned} E_1 &= \{ \text{"≤XS"}, \text{"S"}, \text{"M"} \} \\ E_2 &= \{ \text{"M"}, \text{"L"} \} \\ E_1 \cup E_2 &= \{ \text{"≤XS"}, \text{"S"}, \text{"M"}, \text{"L"} \} \\ E_1 \cap E_2 &= \{ \text{"M"} \} \end{aligned}$$

Discrete Sample Space:

$$\begin{aligned} \mathbb{P}\left(\bigcup_{i=1}^n \{s_i\}\right) &= \mathbb{P}\left(\bigcup_{i=1}^{n-1} (s_i) \cup s_n\right) \\ &= \mathbb{P}\left(\bigcup_{i=1}^{n-1} (s_i)\right) + \mathbb{P}(s_n) \end{aligned}$$

$$= \sum_{i=1}^n \mathbb{P}(\{s_i\})$$

$$= 1$$

point:

interval:

$\Rightarrow$ Probability $\mathbb{P}$: a function maps outcomes to real number

$$\textcircled{1} \quad 0 \leq \mathbb{P}(\{s_i\}) \leq 1, \quad \forall s_i \in S.$$

$$\textcircled{2} \quad \sum_{i=1}^n \mathbb{P}(\{s_i\}) = 1$$

ex). $E_1 = \{1, 2, 3\}, E_2 = \{1, 5, 6\}$

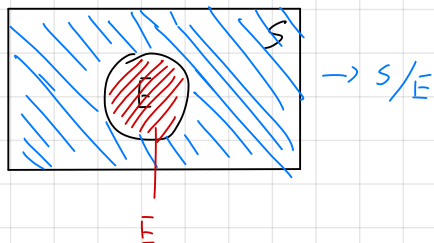
$$\mathbb{P}(E_1) = \frac{1}{2} = \mathbb{P}(E_2) \rightarrow \text{Uniform Probability}$$

Venn Diagram

$$\text{Event } E, \quad \bar{E} = S \setminus E, \quad \mathbb{P}(\bar{E}) = 1 - \mathbb{P}(E).$$

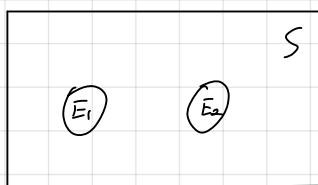
↑
complement of E

$$\mathbb{P}(E) + \mathbb{P}(\bar{E}) = \mathbb{P}(E \cup \bar{E}) = \mathbb{P}(S) = 1.$$



$$\textcircled{1} \mathbb{P}(E_1 \cup E_2) = \mathbb{P}(E_1) + \mathbb{P}(E_2), \text{ if } E_1 \cap E_2 = \emptyset$$

↑
mutually exclusive.



$$\textcircled{2} \mathbb{P}(E_1 \cup E_2) = \mathbb{P}(E_1) + \mathbb{P}(E_2) - \mathbb{P}(E_1 \cap E_2).$$

$$\mathbb{P}(E_1 \cup E_2) = \mathbb{P}(E_1 \setminus (E_1 \cap E_2)) + \mathbb{P}(E_2 \setminus (E_1 \cap E_2)) + \mathbb{P}(E_1 \cap E_2).$$

$$= \mathbb{P}(E_1) - \mathbb{P}(E_1 \cap E_2) + \mathbb{P}(E_2) - \mathbb{P}(E_1 \cap E_2) + \mathbb{P}(E_1 \cap E_2)$$

$$= \mathbb{P}(E_1) + \mathbb{P}(E_2) - \mathbb{P}(E_1 \cap E_2).$$

ex). 0.3 prob. roommates home $\rightarrow A$

0.2 prob. roommates partner home $\rightarrow B$

0.1 prob. both of them home $\rightarrow A \cap B$

Prob no one is at home $\leftarrow$ let C : no one is at home

$$\mathbb{P}(C) = 1 - \mathbb{P}(\bar{C}) = 1 - \mathbb{P}(A \cup B)$$

$$= 1 - (\mathbb{P}(A) + \mathbb{P}(B) - \mathbb{P}(A \cap B))$$

$$= 1 - 0.3 - 0.2 + 0.1$$

$$= 0.6.$$

Sample space = $\{(K \text{ home}, P \text{ home}),$
 $(K \text{ not } h, P h),$
 $(K h, P \text{ not } h),$
 $(K \text{ not } h, P \text{ not } h)\}$

$$\textcircled{3} \quad \mathbb{P}(E_1 \cap E_2) = \mathbb{P}(E_1) \underbrace{\mathbb{P}(E_2 | E_1)}_{\substack{\uparrow \\ \text{prob of } E_2, \text{ given condition on } E_1}} = \mathbb{P}(E_2) \cdot \mathbb{P}(E_2 | E_1)$$

$$E_1, E_2 \rightarrow \mathbb{P}(E_1 \cap E_2) = \mathbb{P}(E_1) \mathbb{P}(E_2)$$

$\uparrow$
independent



$$\mathbb{P}(E_2) = \mathbb{P}(E_2 | E_1).$$

$\textcircled{4}$ Are mutually exclusive event independent?

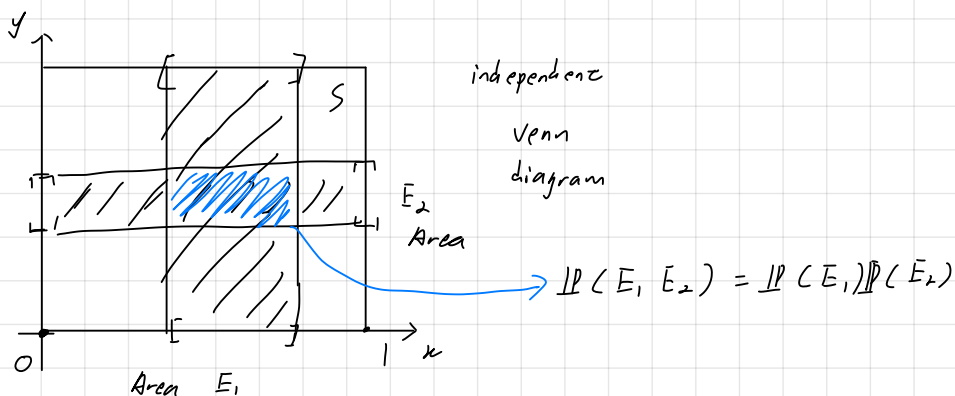
$$\mathbb{P}(E_1 \cap E_2) = \mathbb{P}(\emptyset) = 0$$

$$\mathbb{P}(E_1 | E_2) = 0 \quad \text{if} \quad \mathbb{P}(E_2) \neq 0.$$

$$= \frac{\mathbb{P}(E_1 \cap E_2)}{\mathbb{P}(E_2)} \leftarrow = 0 \quad \leftarrow \neq 0$$

a. if $\mathbb{P}(E_1) \neq 0$ & $\mathbb{P}(E_2) \neq 0$, $E_1 \cap E_2 = \emptyset \Rightarrow E_1, E_2$ not independent

b.



ex) Roll a die 3 times

prob $\{ \text{face 1 never appears} \}$

prob $\{ \text{face 1 appears at least once} \} \leftarrow B_1$

sample space for each roll: $\{1, \dots, 6\}$

$A_1 = \{ \text{roll 1 = face 1} \}$. $\mathbb{P}(A_1) = \frac{1}{6}$

$\bar{A}_1 = \{ \text{roll 1} \neq \text{face 1} \}$. $\mathbb{P}(\bar{A}_1) = \frac{5}{6}$

$$\begin{aligned}\mathbb{P}(\bar{A}_1 \cap \bar{A}_2 \cap \bar{A}_3) &= \mathbb{P}(\bar{A}_1) \cdot \mathbb{P}(\bar{A}_2) \cdot \mathbb{P}(\bar{A}_3) \\ &= \left(\frac{5}{6}\right)^3.\end{aligned}$$

$$\begin{aligned}\mathbb{P}(\bar{B}_1) &= 1 - \mathbb{P}(\bar{A}_1 \cap \bar{A}_2 \cap \bar{A}_3) \\ &= 1 - \left(\frac{5}{6}\right)^3\end{aligned}$$

ex) Two pets Equal prob. $\{ \text{cat, dog} \}$.

Q1 Prob both are dogs, given oldest is dog? (D, D)

Q2 - - - -, at least one of them is dog? $(D, C), (C, D)$
weaker condition (D, D) .

Q1, Event $B = \{ \text{both are dogs} \}$ $\mathbb{P}(B \cap A_1) = \mathbb{P}(B) = \left(\frac{1}{2}\right)^2 = \frac{1}{4}$

Event $A_1 = \{ \text{oldest is dog} \}$ $\mathbb{P}(A_1) = \frac{1}{2}$

$$\mathbb{P}(B|A_1) = \frac{\mathbb{P}(B \cap A_1)}{\mathbb{P}(A_1)} = \frac{\frac{1}{4}}{\frac{1}{2}} = \frac{1}{2}$$

Q2 Event $A_2 = \{ \text{at least one is dog} \}$ $\mathbb{P}(A) = 1 - \mathbb{P}(\bar{A}_2) = \frac{3}{4}$

$\bar{A}_2 = \{ \text{none of them are dogs} \}$ $\mathbb{P}(\bar{A}_2) = \frac{1}{4}$

$$\mathbb{P}(B|A_2) = \frac{\mathbb{P}(B \cap A_2)}{\mathbb{P}(A_2)} = \frac{\frac{1}{4}}{\frac{3}{4}} = \frac{1}{3}$$

$$\mathbb{P}(B \cap A_2) = \mathbb{P}(B) = \frac{1}{4}$$

Condition Probability

$$P(B|A)$$

↑
condition

Law of Total Probability

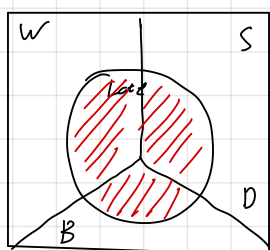
ex)

		Late	on time
Getting to campus	walk :	0.06	0.24
	bike :	0.03	0.07
	drive :	0.36	0.24

Q1 Prob a random is late?

$$P(\text{Late}) = P(\text{Late} \cap \text{walk}) + P(\text{Late} \cap \text{bike}) + P(\text{Late} \cap \text{drive}).$$

$$= P(\text{Late} | \text{walk}) P(\text{walk}) + P(\text{Late} | \text{bike}) P(\text{bike}) + P(\text{Late} | \text{drive}) P(\text{drive}).$$



$$(L \cap W) \cup (L \cap B) \cup (L \cap D) = L \leftarrow \text{disjoint}$$

Property $(A \cap B) \cup (A \cap C) = A \cap (B \cup C)$

$$P(L \cap D) + P(L \cap B) + P(L \cap W) = P(L)$$

Partition $\{E_1, \dots, E_k\}$ is partition of S if

① $P(E_i \cap E_j) = 0$, $\forall i \neq j$. $(E_i \cap E_j = \emptyset, \forall i \neq j)$.

② $\sum_{i=1}^k P(E_i) = 1$ $(\bigcup_{i=1}^k E_i = S)$.

③ $P(A) = \sum_{i=1}^k P(A \cap E_i)$.

$$= \sum_{i=1}^k P(A | E_i) P(E_i)$$

$$Q2 \quad P(\text{Walk} | \text{Late}) = \frac{P(\text{Late} \cap \text{Walk})}{P(\text{Late})} = \frac{0.06}{0.06 + 0.03 + 0.36}$$

$$Q3 \quad P(\text{Late}) = 0.45$$

$$P(\text{Walk}) = 0.3$$

$$P(\text{Late} | \text{Walk}) = 0.2$$

$$P(\text{Walk} | \text{Late}) = \frac{P(\text{Walk} \cap \text{Late})}{P(\text{Late})} = \frac{P(\text{Late} | \text{Walk}) P(\text{Walk})}{P(\text{Late})}$$



Bayes Theorem

$$P(A|B) = \frac{P(B|A) P(A)}{P(B)}$$

Statistical Model



$$- \frac{P(\text{Data} | \text{Parameter}) P(\text{Parameter})}{P(\text{Data})} = P(\text{Parameter} | \text{Data})$$

$\uparrow$
 prior belief $P(\text{Known} | \text{Data 1})$

ex) Product detects Covid-19 w/ 95% prob.

15% prob. turning positive for people w/o covid-19.

10% have Covid-19

Q: Prob. people get covid-19, if they test positive?

T: { test positive }

H: { have covid-19 }

$$P(T) = P(T|H)P(H) + P(T|H^c)P(H^c)$$

$$P(T|H) = 0.95$$

$$P(T|H^c) = 0.15$$

$$P(H) = 0.1$$

$$P(H|T) = \frac{P(T|H) P(H)}{P(T)}$$

$$0.95 \cdot 0.1$$

$$= 0.95 \cdot 0.1 + 0.15 \cdot 0.9$$

$$\approx 0.4$$

Differential Privacy

338 website takes a poll

Republic.

Vem.

49%

51%

- Differential attack.
- Privacy - accuracy trade off
- Guard against privacy leakage

Flip a coin — "head" $\rightarrow$ answer faithfully

"tai!" $\rightarrow$ answer 50% - 50%

Rep.	Dem.
------	------

Q: How safe is this?

Periode "Rep" - 0

"Dem" - 1

Private preference - $X_n \in \{0, 1\}$

Observation $- Y_n \in \{0, 1, \dots, J\}$

private

↓

Question: $\mathbb{P}(X_n = 1 \mid Y_n = 0) = ? \longrightarrow \mathbb{P}(\text{prefer Dem.} \mid \text{Answered "Rep"})$.

$$= \frac{\mathbb{P}(Y_1 = 0 \mid X_1 = 1) \mathbb{P}(X_1 = 1)}{\mathbb{P}(Y_1 = 0)}$$

↑
observed

$$1 \quad \mathbb{P}(Y_1 = 0 \mid X_1 = 0)$$
$$\mathbb{P}(Y_n = 0 \mid X_n = 1) = \mathbb{P}(\text{fairchful}) \mathbb{P}(Y_n = 0 \mid X_n = 1, \text{fairchful}) + \mathbb{P}(\text{unfairchful}) \mathbb{P}(Y_n = 0 \mid X_n = 1, \text{unfairchful}) = \frac{1}{2} \cdot 1 + \frac{1}{2} \cdot \frac{1}{2}$$
$$= \frac{1}{2} \cdot 0 + \frac{1}{2} \cdot \frac{1}{2} = \frac{1}{4}$$

Law of Total Probability

$\underline{P}(X_n = 1) = 0.51$
 $\underline{P}(X_n = 0) = 0.49$ } Assume
 Distribution
 of Y_n

$$\begin{aligned} \mathbb{P}(Y_n=0) &= \mathbb{P}(Y_n=0 | X_n=0) \mathbb{P}(X_n=0) + \\ &\quad \mathbb{P}(Y_n=0 | X_n=1) \mathbb{P}(X_n=1) \end{aligned}$$
$$= \frac{3}{4} (0.49) + \frac{1}{4} (0.51)$$
$$= 0.495$$

$$\frac{\mathbb{P}(Y_n = 0 | X_n = 1) \mathbb{P}(X_n = 1)}{\mathbb{P}(Y_n = 0)} = \frac{\frac{1}{4} - 0.51}{0.495} \approx \frac{1}{4}$$

$$\mathbb{P}(X_n = 0 | Y_n = 0) = 1 - \mathbb{P}(X_n = 1 | Y_n = 0) \approx \frac{3}{4}$$

prob that attacker success.

Q: Can we recover the poll result?

Don't know $\mathbb{P}(X=1)$ but know

$\mathbb{P}(Y=1)$.

Expectation:

$$\frac{1}{n} \sum_{i=1}^n X_i = \frac{1}{n} \left(\sum_{X_i \in D_0} X_i + \sum_{X_i \in D_1} X_i \right)$$

$$n = n_0 + n_1 \quad \uparrow \quad \uparrow$$

$$\begin{matrix} \text{Rep.} & \text{Dem.} \\ D_0 & D_1 \end{matrix} = \frac{n_0}{n} \cdot 0 + \frac{n_1}{n} \cdot 1$$

$$\text{As } n \rightarrow \infty = \mathbb{P}(X=0) \cdot 0 + \mathbb{P}(X=1) \cdot 1$$

$$= \mathbb{P}(X=1)$$

$$S_0, \quad \frac{1}{n} \sum_{i=1}^n Z_i \xrightarrow{n \rightarrow \infty} \sum_{k=1}^K \mathbb{P}(Z = a_k) a_k$$

Expectation.

← PMF
↑ value
pmf
take

$$\text{Similarly, } \frac{1}{n} \sum_{i=1}^n Y_i = \mathbb{P}(Y=0) \cdot 0 + \mathbb{P}(Y=1) \cdot 1$$

$$= \mathbb{P}(Y=1) \quad \leftarrow \text{law of total probability}$$

Expectation of Y

$$= \mathbb{P}(Y=1 | X=0) \mathbb{P}(X=0) + \mathbb{P}(Y=1 | X=1) \mathbb{P}(X=1)$$

$$= \frac{1}{4} \mathbb{P}(X=0) + \frac{3}{4} \mathbb{P}(X=1)$$

$$= \frac{1}{4} (1 - \mathbb{P}(X=1)) + \frac{3}{4} \mathbb{P}(X=1)$$

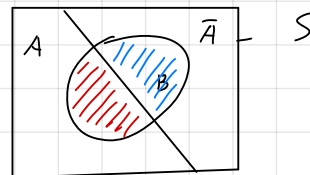
$$= \frac{1}{4} - \frac{1}{4} \mathbb{P}(X=1) + \frac{3}{4} \mathbb{P}(X=1)$$

proof:

$$\mathbb{P}(A|B) = \mathbb{P}(\bar{A}|B)$$

$$\frac{\mathbb{P}(A \cap B)}{\mathbb{P}(B)} = \frac{\mathbb{P}(B)}{\mathbb{P}(B)} - \frac{\mathbb{P}(\bar{A} \cap B)}{\mathbb{P}(B)}$$

$$\frac{\mathbb{P}(A \cap B)}{\mathbb{P}(B)} = \frac{\mathbb{P}(B) - \mathbb{P}(\bar{A} \cap B)}{\mathbb{P}(B)}$$



$$(A \cap B) \cup (\bar{A} \cap B) = B$$

$$\Downarrow$$

$$\mathbb{P}(A \cap B) + \mathbb{P}(\bar{A} \cap B) = \mathbb{P}(B)$$

$$= \frac{1}{4} + \frac{1}{2} \mathbb{P}(X=1).$$

↑
Expectation of X

$\mathbb{P}(Y=1) = \frac{1}{4} + \frac{1}{2} \mathbb{P}(X=1)$ This is how to recover the result!

$$\mathbb{P}(X=1) = 2(\mathbb{P}(Y=1) - \frac{1}{4}).$$

Probability on continuous space

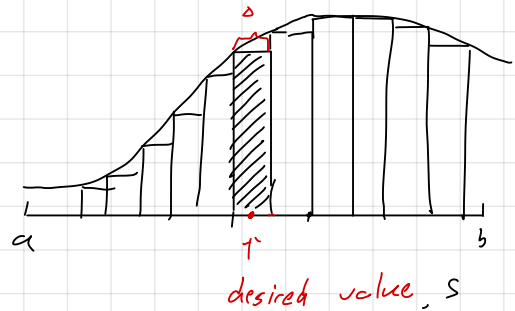
Discrete: $s \in S$, ① $0 \leq \mathbb{P}(s) \leq 1$.

② $\sum_{s \in S} \mathbb{P}(s) = 1$.

$$\lim_{\substack{k \rightarrow \infty \\ \Delta \rightarrow 0}} \sum_{k=1}^k \mathbb{P}(X \in [s_k - \frac{\Delta}{2}, s_k + \frac{\Delta}{2}]) = 1.$$

$$= \lim_{\Delta \rightarrow 0} \sum_{k=1}^k p(s_k) \cdot \Delta$$

$$= \int_a^b p(s) ds.$$



$$\lim_{\Delta \rightarrow 0} \frac{\mathbb{P}(X \in [s - \frac{\Delta}{2}, s + \frac{\Delta}{2}])}{\Delta} = p(s)$$

↑
Probability Density Function

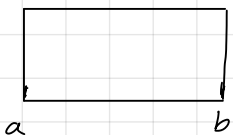
Continuous:

① $p(s) \geq 0$

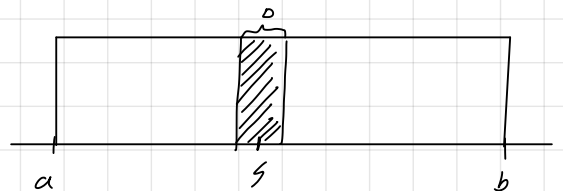
② $\int_{s \in S} p(s) ds = 1$.

Uniform probability case

* Uniform Distribution.



p.d.f. $p(s) = \frac{1}{b-a}$



$$\mathbb{P}(X \in [s - \frac{\Delta}{2}, s + \frac{\Delta}{2}]) = \frac{\Delta}{b-a}$$

• Normal Distribution.

p.d.f. defined on $\mathbb{R}$.

p.d.f. $p(s) = \frac{1}{\sqrt{2\pi} \cdot \sigma} \exp(-\frac{(s-\mu)^2}{2\sigma^2})$

$$\int p(s) ds = 1$$

$$\begin{aligned} &\Leftarrow \int_{\mathbb{R}} \exp(-\frac{(s-\mu)^2}{2\sigma^2}) ds \\ &= \sqrt{2\pi} \cdot \sigma \end{aligned}$$

μ : mean / Expectation of distribution

↓

① Discrete Expectation: $\frac{1}{n} \sum_{i=1}^n x_i \xrightarrow{n \rightarrow \infty} E[X]$

$$E[X] = \sum_{k=1}^K \mathbb{P}(X = s_k) \cdot s_k$$

② Continuous Expectation: $E[X] = \lim_{\Delta \rightarrow 0} \sum_{k=1}^K \mathbb{P}(X \in [s_k - \frac{\Delta}{2}, s_k + \frac{\Delta}{2}]) \cdot s_k$

$$= \lim_{\Delta \rightarrow 0} \sum_{k=1}^K p(s_k) \Delta \cdot s_k$$

$$= \int_{s \in S} p(s) \cdot s \, ds.$$

For normal distribution: $E[X] = \int_{\mathbb{R}} p_N(s) \cdot s \, ds = \mu$

σ : STD, σ^2 : Variance

$$\sigma^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2$$

$$\downarrow n \rightarrow \infty$$
$$\int_{s \in S} p(s) \cdot (s - \mu)^2 \, ds = E[(X - \mu)^2]$$

↑

$$\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 \xrightarrow{n \rightarrow \infty} \sum_{k=1}^K \mathbb{P}(X = s_k) \cdot (s_k - \mu)^2$$

$$\frac{1}{\sqrt{2\pi} \cdot \sigma} \int_{-\infty}^{\infty} \exp(-(s-\mu)^2 / 2\sigma^2) (s-\mu)^2 \, ds.$$

$$\text{let } y = s - \mu$$

$$\frac{1}{\sqrt{2\pi} \cdot \sigma} \int_{-\infty}^{\infty} \exp\left(-\frac{y^2}{2\sigma^2}\right) (y^2) \, dy.$$

do integral by parts twice

$$\exp(-y^2/2\sigma^2) = \exp(-y^2/2\sigma^2) \cdot \left(\frac{-2y}{2\sigma^2}\right) dy$$

$$E[f(x)]$$
$$= \sum_{k=1}^K \mathbb{P}(X = s_k) \cdot f(s_k).$$

$$\begin{aligned}
& \frac{1}{\sqrt{2\pi} \cdot \sigma} \int_{-\infty}^{\infty} \exp\left(\frac{-y^2}{2\sigma^2}\right) (y^2) dy \\
&= \frac{1}{\sqrt{2\pi} \cdot \sigma} \int_{-\infty}^{\infty} \left(-\frac{\sigma^2}{y}\right) \cdot d\exp\left(-y^2/2\sigma^2\right) \\
&= -\frac{1}{\sqrt{2\pi} \cdot \sigma} \int_{-\infty}^{\infty} \exp\left(\frac{-y^2}{\sigma^2}\right) \cdot (-\sigma^2) dy \\
&= \sigma^2 \cdot \frac{1}{\sqrt{2\pi} \cdot \sigma} \underbrace{\int_{-\infty}^{\infty} \exp(-y^2/\sigma^2) dy}_{\sqrt{2\pi} \cdot \sigma} \\
&= \sigma^2
\end{aligned}$$

$$X \sim N(\mu, \sigma^2) \quad p(x) = \frac{1}{\sqrt{2\pi} \cdot \sigma} \exp(-(x-\mu)^2/2\sigma^2)$$

$$E[X] = \int_{-\infty}^{\infty} p(x) \cdot x dx = \mu$$

$$\text{Var}(X) = E[(X-\mu)^2] = \int_{-\infty}^{\infty} p(x) (x-\mu)^2 dx = \sigma^2$$

$$E[f(x)] = \int_{-\infty}^{\infty} p(x) f(x) dx$$

$$E[X+Y] = \iint_{\mathbb{R}^2} p_{(X,Y)}(x,y) \cdot (x+y) dx dy$$

$$\begin{aligned}
& \int_{-\infty}^{\infty} p_{(X,Y)}(x,y) dy \\
&= p(x)
\end{aligned}$$



$$\begin{aligned}
& \iint_{\mathbb{R}^2} p_{(X,Y)}(x,y) \cdot x dx dy + \iint_{\mathbb{R}^2} p_{(X,Y)}(x,y) \cdot y dx dy \\
&= \int_{\mathbb{R}} x \int_{\mathbb{R}} p_{(X,Y)}(x,y) dy dx + \int_{\mathbb{R}} y \int_{\mathbb{R}} p_{(X,Y)}(x,y) dx dy \\
&= \int_{\mathbb{R}} x p_X(x) dx + \int_{\mathbb{R}} y p_Y(y) dy \\
&= E[X] + E[Y]
\end{aligned}$$

$$p_{X,Y}(x,y) = \lim_{\Delta \rightarrow 0} \frac{\mathbb{P}(X \in [x-\frac{\Delta}{2}, x+\frac{\Delta}{2}] \cap Y \in [y-\frac{\Delta}{2}, y+\frac{\Delta}{2}])}{\Delta^2}$$

$$\int_{-\infty}^{\infty} p_{X,Y}(x,y) dy = \sum_{\Delta=1}^{\infty} \frac{\mathbb{P}(X \in [x-\frac{\Delta}{2}, x+\frac{\Delta}{2}] \cap Y \in [y-\frac{\Delta}{2}, y+\frac{\Delta}{2}])}{\Delta^2} \cdot \Delta$$

$$= \frac{\mathbb{P}(X \in [x-\frac{\Delta}{2}, x+\frac{\Delta}{2}] \cap Y \in \mathbb{R})}{\Delta}$$

$$= \frac{\mathbb{P}(X \in [x-\frac{\Delta}{2}, x+\frac{\Delta}{2}])}{\Delta} = p_X(x)$$

Independence & Expectation

$$\mathbb{P}(A \cap B) = \mathbb{P}(A) \mathbb{P}(B)$$

$$p_{(X,Y)}(x,y) = p_X(x) p_Y(y)$$

b/c

$$p_{X,Y}(x,y) = \lim_{\Delta \rightarrow 0} \frac{\mathbb{P}(X \in [x - \frac{\Delta}{2}, x + \frac{\Delta}{2}]) \cdot \mathbb{P}(Y \in [y - \frac{\Delta}{2}, y + \frac{\Delta}{2}])}{\Delta^2} \quad \text{when independence}$$

$$= \lim_{\Delta \rightarrow 0} \underbrace{\frac{1}{\Delta} \mathbb{P}(X \in [x - \frac{\Delta}{2}, x + \frac{\Delta}{2}])}_{\downarrow p_X(x)} \underbrace{\frac{1}{\Delta} \mathbb{P}(Y \in [y - \frac{\Delta}{2}, y + \frac{\Delta}{2}])}_{\downarrow p_Y(y)}$$

$$E[XY] = \int \int_{\mathbb{R}} xy p_X(x) p_Y(y) dx dy \quad \text{for independence.}$$

$$= \int x p_X(x) \int y p_Y(y) dx dy$$

$$= \int x p_X(x) E[Y] dx$$

$$= E[X] E[Y]$$

$$X \sim N(\mu_1, \sigma_1^2) \Rightarrow X+Y \sim N(\mu_1 + \mu_2, \sigma_1^2 + \sigma_2^2) \quad \text{when independence.}$$

$$Y \sim N(\mu_2, \sigma_2^2)$$

b/c

$$E[X+Y] = E[X] + E[Y]$$

$$\text{Var}(X+Y) = E[(X+Y - E[X+Y])^2]$$

$$= E[(X - E[X]) + (Y - E[Y])^2]$$

$$= E[(X - E[X])^2 + (Y - E[Y])^2 + 2(X - E[X])(Y - E[Y])]$$

$$= E[(X - E[X])^2] + E[(Y - E[Y])^2] + E[2(X - E[X])(Y - E[Y])]$$

$$= \text{Var}(X) + \text{Var}(Y) + 2 \underbrace{E[(X - E[X])(Y - E[Y])]}_{\downarrow 0}$$

$$\downarrow \sigma^2$$

$$\downarrow \sigma^2$$

$$= E[X - E[X]] E[Y - E[Y]]$$

$$= E[X] - E[E[X]] \dots$$

$$= E[X] - E[X] \dots$$

$$= 0$$

$$\bullet X + X = 2X \sim N(2\mu, \underline{4\sigma^2}), \quad X + Y \sim N(2\mu, \underline{2\sigma^2})$$

$$E[2X] = 2E[X] = 2\mu$$

↑
independence

due to
cancellation

$$E[(2X - E[2X])^2] = 4E[(X - E[X])^2] \\ = 4\sigma^2$$

Empirical Risk Minimization

Go back to differential privacy:

$$\bullet \text{loss}(h, x_i)$$

Mechanism: w/ x probability faithfully

$$\bullet R(H; D) = \frac{1}{n} \sum_{i=1}^n \text{loss}(h, x_i)$$

w/ $1-x$ probability unfaithfully
random

Can attacker deduce x , based on data?

Assume they know: $\mathbb{P}(\text{"Dem"}) = \frac{1}{4}$.

Observe $\{\text{"D"}, \text{"R"}, \text{"K"}, \text{"D"}, \text{"R"}, \text{"R"}\} = S$.

If we know $\mathbb{P}(X \in [d - \frac{\Delta}{2}, d + \frac{\Delta}{2}] \mid \text{Data} = S)$.

$$a^* = \arg \min_a \mathbb{P}(X \in [d - \frac{\Delta}{2}, d + \frac{\Delta}{2}] \mid \text{Data} = S) \\ d^* = \arg \min_a P_X(d \mid \text{Data} = S)$$

Bayes's
↓

$$= \frac{\mathbb{P}(\text{Data} = S \mid X \in \dots) \mathbb{P}(X \in \dots)}{\mathbb{P}(\text{Data} = S)}$$

$$\text{For small } \Delta = \frac{\mathbb{P}(\text{Data} = S \mid X = d) P_X(d)}{\mathbb{P}(\text{Data} = S)}$$

$$\propto \mathbb{P}(\text{Data} = S \mid X = d) P_X(d)$$

↑
likelihood

↑
prior probability

Probability Theory for ERM

DP problem with x prob, reveal true preference
($1-x$) prob, uniform random.

$$\mathbb{P}(\text{"Dem"}) = \frac{1}{4}, \quad \mathbb{P}(\text{"Rep"}) = \frac{3}{4}. \quad \text{Observe } \{\text{"D"}, \text{"R"}, \text{"K"}, \text{"D"}, \text{"R"}, \text{"R"}\} = S.$$

$$\mathbb{P}(x \in [d - \frac{\Delta}{2}, d + \frac{\Delta}{2}] | \text{Data} = s) = p_x(x | \text{Data} = s) \quad \uparrow$$

pick $x^* = \underset{d}{\operatorname{argmax}} \mathbb{P}(x \in [d - \frac{\Delta}{2}, d + \frac{\Delta}{2}] | \text{Data} = s).$

Posterior

Bayes's Thm: $\mathbb{P}(x \in [d - \frac{\Delta}{2}, d + \frac{\Delta}{2}] | \text{Data} = s).$

$$= \mathbb{P}(\text{Data} = s | x \in [d - \frac{\Delta}{2}, d + \frac{\Delta}{2}]) \mathbb{P}(x \in [d - \frac{\Delta}{2}, d + \frac{\Delta}{2}]) / \mathbb{P}(\text{Data} = s).$$

$$\propto \mathbb{P}(\text{Data} = s | x \in [d - \frac{\Delta}{2}, d + \frac{\Delta}{2}]) \mathbb{P}(x \in [d - \frac{\Delta}{2}, d + \frac{\Delta}{2}]).$$

small Δ

$$\propto \mathbb{P}(\text{Data} = s | x = d) p_x(d)$$

$\uparrow$
Likelihood

$\uparrow$
prior probability / belief

Maximum Likelihood Estimate

$$\mathbb{P}(\text{Data} = s | x = d) = \prod_{i=1}^6 \mathbb{P}(\text{Data}(i) = s(i) | x = d).$$

$\Downarrow$

$$\text{Data}(1) = s(1) \wedge \text{Data}(2) = s(2) \dots$$

$$\begin{aligned} \mathbb{P}(\text{Data} = "D" | x = d) &= d \cdot \frac{1}{4} + (1-d) \cdot \frac{1}{2} \\ &= \frac{1}{2} - \frac{1}{4}d \end{aligned}$$

$$\begin{aligned} \mathbb{P}(\text{Data} = "K" | x = d) &= 1 - \mathbb{P}(\text{Data} = "D" | x = d) \\ &= \frac{1}{2} + \frac{1}{4}d \end{aligned}$$

$\uparrow$ independence

Empirical Risk

$\downarrow$

$$\text{Risk}(d | s) = -\ln \mathbb{P}(\text{Data} = s | x = d)$$

$$= -2 \ln(\frac{1}{2} - \frac{1}{4}d) - 4 \ln(\frac{1}{2} + \frac{1}{4}d).$$

$$\mathbb{P}(\text{Data}(i+1) | \text{Data}(i), x = d)$$

Not independent case

$$\text{so } \frac{d}{d d} R(d | s) = -2 \frac{-\frac{1}{4}}{\frac{1}{2} - \frac{1}{4}d} - 4 \frac{\frac{1}{4}}{\frac{1}{2} + \frac{1}{4}d} = 0$$

$\Downarrow$

$$\frac{1}{2} \cdot \frac{1}{\frac{1}{2} - \frac{1}{4}a} = \frac{1}{\frac{1}{2} + \frac{1}{4}a}$$

$$\frac{1}{2}(\frac{1}{2} + \frac{1}{4}a) = \frac{1}{2} - \frac{1}{4}a$$

$$\frac{1}{4} + \frac{1}{8}a = \frac{1}{2} - \frac{1}{4}a$$

$$\frac{3}{8}a = \frac{1}{4}$$

$$a = \frac{2}{3}$$

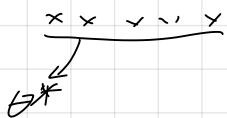
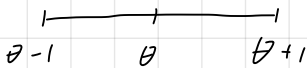
Q: What if the calculated a is negative?

||

prior belief

$$p_x(a) = \begin{cases} 1, & 0 \leq a \leq 1 \\ 0, & \text{else} \end{cases}$$

Relating to groupwork



$P(\theta | D)$: Posterior

$$P(D | \theta) : \text{likelihood} = \prod_{i=1}^n P(x_i | \theta)$$

$$x_i \sim \text{Unif}[\theta-1, \theta+1]$$

$$P_{x_i}(n) = \begin{cases} \frac{1}{2}, & x_i \in [\theta-1, \theta+1] \\ 0, & \text{else} \end{cases}$$

$$\text{if } x_1, \dots, x_n \in [\theta^*-1, \theta^*+1]$$

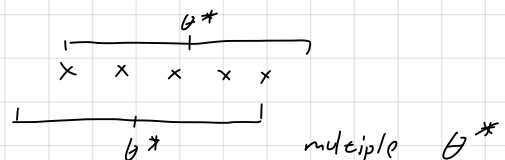
$\Downarrow$

$$P(D | \theta) > 0$$

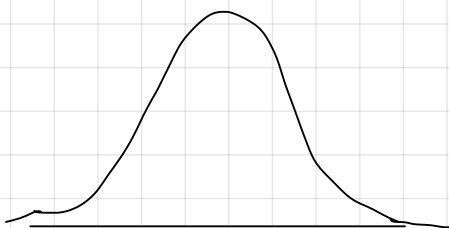
$$\text{if } \exists x_i \notin [\theta^*-1, \theta^*+1],$$

$\Downarrow$

$$P(D | \theta) = 0$$



Central Limit Theorem :



$$P(\theta | D) \propto \prod_{i=1}^n P(x_i | \theta) \cdot P(\theta) \rightarrow \text{Normal Distribution}$$

Bernstein von Mises

Q: Why not Model data as normal R.V.?

$$P(x_i | \theta): x_i \sim N(\theta, \sigma^2).$$

$$\begin{aligned} \text{Then } P(D | \theta) &= \prod_{i=1}^n P(x_i | \theta) \quad \text{likelihood} \\ &= \prod_{i=1}^n \frac{1}{\sqrt{2\pi} \cdot \sigma} e^{-(x_i - \theta)^2 / 2\sigma^2} \\ &\propto \prod_{i=1}^n e^{-(x_i - \theta)^2 / 2\sigma^2} \\ &= e^{-\sum_{i=1}^n (x_i - \theta)^2 / 2\sigma^2} \\ &= e^{-\frac{1}{2\sigma^2} \sum_{i=1}^n (x_i - \theta)^2} \end{aligned}$$

Empirical risk

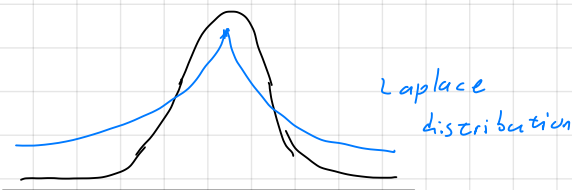
$$\begin{aligned} R(\theta, D) &= -\ln P(D | \theta) \\ &= \frac{1}{2\sigma^2} \sum_{i=1}^n (x_i - \theta)^2 \end{aligned}$$

Empirical risk square d loss

$$\theta^* = \arg \max_{\theta} P(D | \theta) = \arg \min_{\theta} R(\theta, D) = \frac{1}{n} \sum_{i=1}^n x_i$$

With finite data, how to model outliers?

Normal distribution



$$P(x_i | \theta) \propto \exp\left(-\frac{(x_i - \theta)^2}{2\sigma^2}\right) \quad \text{value decreases faster}$$

$$P(x_i | \theta) = \frac{1}{2} \exp(-|x_i - \theta|)$$

$$P(D | \theta) = \prod_{i=1}^n \exp(-|x_i - \theta|) \propto \exp\left(-\sum_{i=1}^n |x_i - \theta|\right)$$

$$R(\theta | D) = \sum_{i=1}^n |x_i - \theta| \quad \text{Absolute loss}$$

Regression

$$y_i \approx H(x_i, \theta)$$

Linear regression

$$y_i \simeq \langle \vec{x}_i, \vec{w} \rangle$$

$$y_i \sim N(\langle \vec{x}_i, \vec{w} \rangle, \sigma^2) \Rightarrow P(y_i | \vec{x}_i, \vec{w})$$

$$P(D | \vec{w}) = \prod_{i=1}^n P((x, y) | \vec{w})$$

$$P((x_i, y_i) | \vec{w}) = P(y_i | \vec{x}_i, \vec{w}) \underbrace{P(\vec{x}_i | \vec{w})}_{P(\vec{x}_i)}$$

$$\propto P(y_i | \vec{x}_i, \vec{w}) P(\vec{x}_i)$$

$$\propto \exp\left(- (y_i - \langle \vec{x}_i, \vec{w} \rangle)^2 / 2 \sigma^2\right).$$

Risk = least square

P.3

A: J

B: J & K A

C: Heart

$$P(A \cap B | C) = \frac{1}{12}$$

$$P(A | C) P(B | C)$$

$$= \frac{1}{12} \cdot \frac{4}{12} = \frac{4}{24} = \frac{1}{6}$$

$$P(A | B \cap C) = \frac{4}{8} = \frac{1}{2}$$

$$P(A | C) = \frac{1}{2}$$

$$P(A | B \cap C) = \frac{1}{4}$$

$$P(A | C) = \frac{1}{12}$$

A: heart

B: face

C: red
condition

$$P(A \cap B | C) = \frac{4}{24} = \frac{1}{6}$$

$$P(A | C) P(B | C) = \frac{12}{24} \cdot \frac{8}{24} = \frac{1}{2} \cdot \frac{1}{3} = \frac{1}{6}$$

$$P(B | C)$$