

*Announcement — Gradescope submission.

Code submission. .PY

[HWO-A ← .pdf
HWO-B ← .pdf
HWO-code ← .py

A	B
✓	✓
	✓
✓	✓

Lecture 3

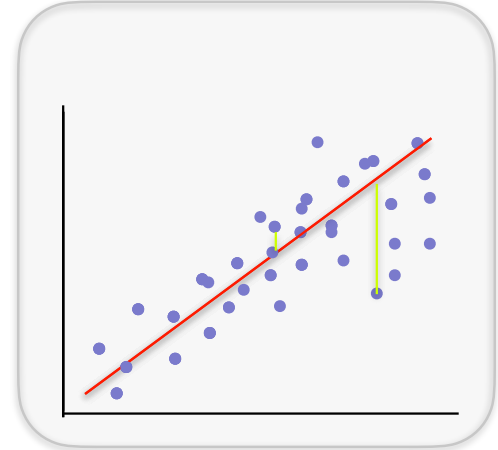


The regression problem in matrix notation

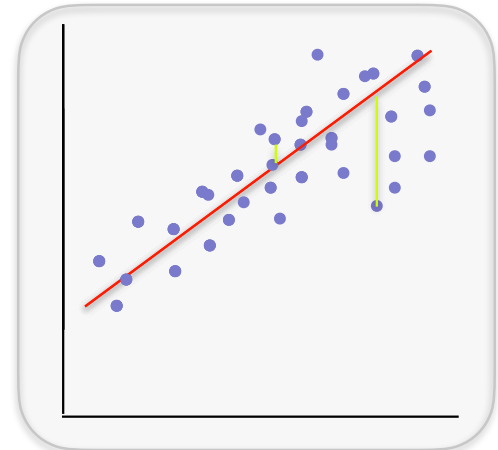
Linear model: $y_i = x_i^T w + \epsilon_i$

Least squares solution:

$$\begin{aligned} \mathbb{R}^d \ni \hat{w}_{LS} &= \arg \min_w ||\mathbf{y} - \mathbf{X}w||_2^2 \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} \end{aligned}$$



What about an offset
(a.k.a intercept)?

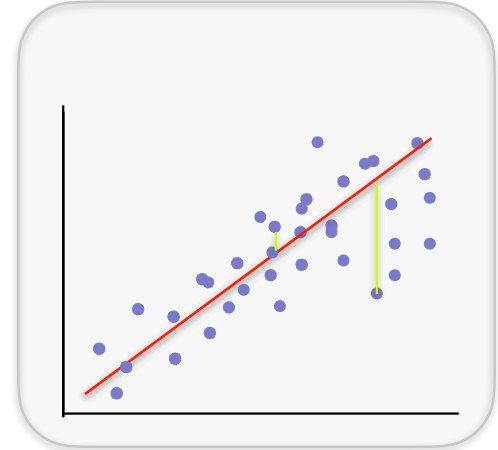


The regression problem in matrix notation

Linear model: $y_i = x_i^T w + \epsilon_i$

Least squares solution:

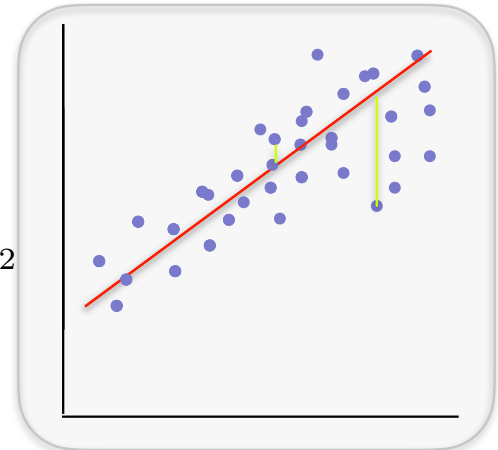
$$\begin{aligned}\hat{w}_{LS} &= \arg \min_w \|\mathbf{y} - \mathbf{X}w\|_2^2 \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}\end{aligned}$$



Affine model: $y_i = x_i^T w + \underset{\substack{\mathbb{R} \\ \mathbb{P}}}{b} + \epsilon_i$

Least squares solution:

$$\begin{aligned}\hat{w}_{LS}, \hat{b}_{LS} &= \arg \min_{w,b} \sum_{i=1}^n (y_i - (x_i^T w + b))^2 \\ &= \arg \min_{w,b} \|\mathbf{y} - (\mathbf{X}w + \mathbf{1}b)\|_2^2\end{aligned}$$



Dealing with an offset

$$\hat{w}_{LS}, \hat{b}_{LS} = \arg \min_{\substack{\mathbb{R}^d \ni w, b \\ \mathbb{R}}} \underbrace{\|y - (Xw + \mathbf{1}b)\|_2^2}_{\substack{\text{Vector } \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix} \left\}^n \\ \text{Quad fun. } w, b.}} \leftarrow \text{Quad fun. } w, b.$$

Set gradient w.r.t. w and b to zero to find the minima:

$$\underbrace{\nabla_w}_{\substack{\mathbb{R}^d \\ \mathbb{R}^d}} L(w, b) = 2 \cdot X^T (y - (Xw + \mathbf{1}b)) \Big|_{w, b=0}$$

$$\underbrace{d}_{n} \underbrace{\boxed{X^T}}_n \left(\underbrace{\boxed{y}}_n - \left(\underbrace{\boxed{X}}_n \underbrace{\boxed{w}}_n + \underbrace{\boxed{\mathbf{1}}}_n \underbrace{\boxed{b}}_n \right) \right)$$

$$\underbrace{\nabla_b}_{\mathbb{R}} L(w, b) = 2 \cdot \mathbf{1}^T (y - (Xw + \mathbf{1}b)) \Big|_{w, b=0}$$

$$\uparrow \text{Fact: } \nabla_w ((Aw + b)^T (Aw + b)) = 2A^T(Aw + b)$$

$$(aw + b)^2 = 2a(aw + b) \quad \left| \begin{array}{l} \nabla_w (w^T w) = 2w, \\ \nabla_w (w^T A^T A w) = \end{array} \right. \rightarrow$$

$$X^T y = X^T X w + X^T \mathbf{1} b$$

$$\mathbf{1}^T y = \mathbf{1}^T X w + \mathbf{1}^T \mathbf{1} b$$

$$\underbrace{\text{Linear in } \substack{\hat{w}, \hat{b} \\ \mathbb{R}^d \mathbb{R}^1}}_{\substack{\text{d+1 variables} \\ \text{d+1 equations}}}$$

d+1 variables
d+1 equations

Dealing with an offset

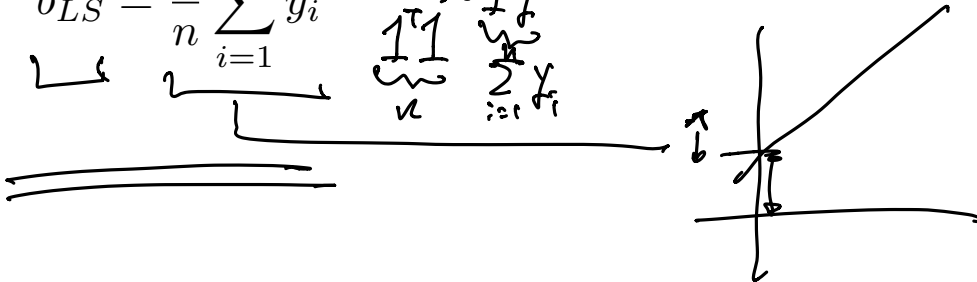
$$\hat{w}_{LS}, \hat{b}_{LS} = \arg \min_{w, b} \|\mathbf{y} - (\mathbf{X}w + \mathbf{1}b)\|_2^2$$

$$\begin{bmatrix} \sum_{i=1}^n x_{i,1} \\ \vdots \\ \sum_{i=1}^n x_{i,d} \end{bmatrix} = \begin{bmatrix} 0 \\ \vdots \\ 0 \end{bmatrix} \quad \begin{matrix} \mathbf{X}^T \mathbf{X} \hat{w}_{LS} + \hat{b}_{LS} \mathbf{X}^T \mathbf{1} = \mathbf{X}^T \mathbf{y} \\ \mathbf{1}^T \mathbf{X} \hat{w}_{LS} + \hat{b}_{LS} \mathbf{1}^T \mathbf{1} = \mathbf{1}^T \mathbf{y} \end{matrix} \rightarrow \begin{matrix} \hat{b} \cdot (\mathbf{1}^T \mathbf{1}) = \mathbf{1}^T \mathbf{y} \\ \hat{b} \cdot n = \sum_{i=1}^n y_i \\ \hat{b} = \frac{1}{n} \sum_{i=1}^n y_i \end{matrix}$$

If $\mathbf{X}^T \mathbf{1} = 0$, if the features have zero mean,

$$\hat{w}_{LS} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y}$$

$$\hat{b}_{LS} = \frac{1}{n} \sum_{i=1}^n y_i = \frac{1}{n} \cdot \frac{\mathbf{1}^T \mathbf{y}}{\mathbf{1}^T \mathbf{1}} = \frac{1}{n} \cdot \frac{\sum_{i=1}^n y_i}{\sum_{i=1}^n 1}$$



Dealing with an offset

$$\mu_j \triangleq \frac{1}{n} \sum_{i=1}^n X_{i,j}$$

μ_1 : average μ_2 : sum
offset. $\neq$ bias/bias.

$$\hat{w}_{LS}, \hat{b}_{LS} = \arg \min_{w, b} \underbrace{\|y - (Xw + \mathbf{1}b)\|_2^2}_{\mathcal{L}(w, b)}$$

$$\begin{aligned} \mathbf{X}^T \mathbf{X} \hat{w}_{LS} + \hat{b}_{LS} \mathbf{X}^T \mathbf{1} &= \mathbf{X}^T \mathbf{y} \\ \mathbf{1}^T \mathbf{X} \hat{w}_{LS} + \hat{b}_{LS} \mathbf{1}^T \mathbf{1} &= \mathbf{1}^T \mathbf{y} \end{aligned} \quad C = \mu^T w + b$$

$$\mu = \frac{1}{n} \mathbf{X}^T \mathbf{1} = \left[\frac{1}{n} \sum_{i=1}^n X_{i,1} \right]$$

If $\mathbf{X}^T \mathbf{1} = 0$,

$$\hat{w}_{LS} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

$$\hat{b}_{LS} = \frac{1}{n} \sum_{i=1}^n y_i$$

$\tilde{\mathbf{X}}, C$: new

In general, when $\mathbf{X}^T \mathbf{1} \neq 0$,

$$\begin{aligned} \mathcal{L}(w, b) &= \|y - (\underbrace{X - \mathbf{1}\mu^T}_{+} w - \underbrace{\mathbf{1}\mu^T w}_{-} - \mathbf{1}b)\|_2^2 \\ &= \|y - \underbrace{\tilde{X} \cdot w}_{\text{centered}} - \mathbf{1} \cdot C\|_2^2, \end{aligned}$$

$$\hat{w}_{LS} = (\tilde{X}^T \tilde{X})^{-1} \tilde{X}^T y$$

$$\begin{aligned} \hat{C}_{LS} &= \frac{1}{n} \sum_{i=1}^n y_i \implies \hat{b}_{LS} = \frac{1}{n} \sum_{i=1}^n y_i - \mu^T \hat{w}_{LS} \\ &\parallel \\ \hat{b}_{LS} + \mu^T \hat{w}_{LS} \end{aligned}$$

Dealing with an offset

$$\hat{w}_{LS}, \hat{b}_{LS} = \arg \min_{w, b} \|\mathbf{y} - (\mathbf{X}w + \mathbf{1}b)\|_2^2$$

$$\mathbf{X}^T \mathbf{X} \hat{w}_{LS} + \hat{b}_{LS} \mathbf{X}^T \mathbf{1} = \mathbf{X}^T \mathbf{y}$$

$$\mathbf{1}^T \mathbf{X} \hat{w}_{LS} + \hat{b}_{LS} \mathbf{1}^T \mathbf{1} = \mathbf{1}^T \mathbf{y}$$

If $\mathbf{X}^T \mathbf{1} = 0$,

$$\hat{w}_{LS} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

$$\hat{b}_{LS} = \frac{1}{n} \sum_{i=1}^n y_i$$

In general, when $\mathbf{X}^T \mathbf{1} \neq 0$,

$$\mu = \frac{1}{n} \mathbf{X}^T \mathbf{1}$$

$$\tilde{\mathbf{X}} = \mathbf{X} - \mathbf{1}\mu^T$$

$$\hat{w}_{LS} = (\tilde{\mathbf{X}}^T \tilde{\mathbf{X}})^{-1} \tilde{\mathbf{X}}^T \mathbf{y}$$

$$\hat{b}_{LS} = \frac{1}{n} \sum_{i=1}^n y_i - \mu^T \hat{w}_{LS}$$

Process

Decide on a **model**: $y_i = x_i^T w + b + \epsilon_i$

Choose a loss function - least squares

Pick the function which minimizes loss on data

$$\hat{w}_{LS}, \hat{b}_{LS} = \arg \min_{w, b} \sum_{i=1}^n (y_i - (x_i^T w + b))^2$$

Use function to make prediction on new examples

$$\hat{y}_{\text{new}} = x_{\text{new}}^T \hat{w}_{LS} + \hat{b}_{LS}$$

Another way of dealing with an offset

$$\hat{w}_{LS}, \hat{b}_{LS} = \arg \min_{w, b} \| \mathbf{y} - (\mathbf{X}w + \mathbf{1}b) \|_2^2$$

reparametrize the problem as $\bar{\mathbf{X}} = [\mathbf{X}, \mathbf{1}]$ and $\bar{w} = \begin{bmatrix} w \\ b \end{bmatrix}$

$$\bar{\mathbf{X}} \bar{w} = \begin{matrix} \hat{\mathbf{X}} & \mathbf{1} \\ \begin{matrix} x_1 \\ \vdots \\ x_n \end{matrix} & \begin{matrix} 1 \\ \vdots \\ 1 \end{matrix} \end{matrix} \cdot \begin{matrix} \bar{w} \\ \left\{ \begin{matrix} w \\ \vdots \\ b \end{matrix} \right\}_1 \end{matrix} = \mathbf{X} \cdot w + \mathbf{1} \cdot b$$

$$\hat{\bar{w}} = \arg \min_{\bar{w}} \| \mathbf{y} - \bar{\mathbf{X}} \bar{w} \|_2^2 \leftarrow (\bar{\mathbf{X}}^T \bar{\mathbf{X}})^{-1} \bar{\mathbf{X}}^T \mathbf{y}$$

$$d+1 \left\{ \begin{bmatrix} \vdots \\ 1 \end{bmatrix} \right\} \begin{matrix} \hat{w} \\ \hat{b} \end{matrix}$$

Why is least squares a good loss function?

$$\begin{aligned}\hat{w}_{LS} &= \arg \min_w ||\mathbf{y} - \mathbf{X}w||_2^2 \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}\end{aligned}$$

Consider $y_i = x_i^T w + \epsilon_i$ where $\epsilon_i \stackrel{i.i.d.}{\sim} \mathcal{N}(0, \sigma^2)$ probabilistic
Generative
Model

$$\Rightarrow y_i \sim \mathcal{N}(x_i^T w, \sigma^2)$$

$$\Rightarrow P(y_i; x_i, w, \sigma) = \frac{1}{\sqrt{2\pi}\sigma} \cdot e^{-\frac{(x_i^T w - y_i)^2}{2\sigma^2}}$$

Why is least squares a good loss function?

Maximum Likelihood Estimator:

$$\begin{aligned}\hat{w}_{\text{MLE}} &= \arg \max_w \log P(\{y_i\}_{i=1}^n; \{x_i\}_{i=1}^n, w, \sigma) \\ &= \arg \max_w -n \log(\sigma \sqrt{2\pi}) + \sum_{i=1}^n -\frac{(y_i - x_i^T w)^2}{2\sigma^2} \\ &= \arg \max_w \sum_{i=1}^n -(y_i - x_i^T w)^2 \\ &= \arg \min_w \sum_{i=1}^n (y_i - x_i^T w)^2\end{aligned}$$

Why is least squares a good loss function?

Maximum Likelihood Estimator:

$$\begin{aligned}\hat{w}_{\text{MLE}} &= \arg \max_w \log P(\{y_i\}_{i=1}^n; \{x_i\}_{i=1}^n, w, \sigma) \\ &= \arg \max_w -n \log(\sigma \sqrt{2\pi}) + \sum_{i=1}^n -\frac{(y_i - x_i^T w)^2}{2\sigma^2} \\ &= \arg \min_w \sum_{i=1}^n (y_i - x_i^T w)^2\end{aligned}$$

Recall: $\hat{w}_{LS} = \arg \min_w \sum_{i=1}^n (y_i - x_i^T w)^2$

$$\hat{w}_{LS} = \hat{w}_{MLE} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y}$$

Recap of linear regression

Data $\{(x_i, y_i)\}_{i=1}^n$

Minimize the loss (Empirical Risk Minimization)

Choose a loss

e.g., $(y_i - x_i^T w)^2$

Solve $\hat{w}_{\text{LS}} = \arg \min_w \sum_{i=1}^n (y_i - x_i^T w)^2$

Maximize the likelihood (MLE)

Choose a Hypothesis class

e.g., $y_i = x_i^T w + \epsilon_i, \quad \epsilon_i \sim \mathcal{N}(0, \sigma^2)$

Maximize the likelihood,

$$\hat{w}_{\text{MLE}} = \arg \max_w \left\{ -n \log(\sigma \sqrt{2\pi}) - \sum_{i=1}^n \frac{(y_i - x_i^T w)^2}{2\sigma^2} \right\}$$

Analysis of **Error** under additive Gaussian noise

$$\text{if } y_i = x_i^T w + \epsilon_i \quad \text{and} \quad \epsilon_i \stackrel{i.i.d.}{\sim} \mathcal{N}(0, \sigma^2) \quad \mathbf{Y} = \mathbf{X}w + \epsilon$$

$$\begin{aligned} \hat{w}_{MLE} &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y} \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T (\mathbf{X}w + \epsilon) \\ &= w + (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \epsilon \end{aligned}$$

Handwritten notes:
- An arrow points from ϵ in the second line to a bracket labeled n , with $\mathcal{N}(0, \sigma^2)$ written above it.
- A circle around ϵ in the third line has an arrow pointing to the word "Random".

Maximum Likelihood Estimator is unbiased:

$$\begin{aligned} \text{Bias}(\hat{w}_{MLE}) &= \mathbb{E}[\hat{w}_{MLE}] - w \\ &\stackrel{11}{=} \mathbb{E}[w + (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \epsilon] \\ &\stackrel{11}{=} w + (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \underbrace{\mathbb{E}[\epsilon]}_{\mathbf{0}} = w \end{aligned}$$

Handwritten notes:
- A bracket labeled n is under ϵ in the second line.
- A bracket labeled 0 is under $\mathbb{E}[\epsilon]$ in the third line.

Analysis of **Error** under additive Gaussian noise

if $y_i = x_i^T w + \epsilon_i$ and $\epsilon_i \stackrel{i.i.d.}{\sim} \mathcal{N}(0, \sigma^2)$ $\mathbf{Y} = \mathbf{X}w + \epsilon$

$$\begin{aligned}\hat{w}_{MLE} &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y} \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T (\mathbf{X}w + \epsilon) \\ &= w + (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \epsilon\end{aligned}$$

Covariance is: $\mathbb{E}[(\hat{w} - w)(\hat{w} - w)^T]$

$$\begin{aligned}&= \mathbb{E}[(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \epsilon \cdot \epsilon^T \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1}] \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \underbrace{\mathbb{E}[\epsilon \epsilon^T]}_{\sigma^2 \mathbf{I}} \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1} = \sigma^2 \underbrace{(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1}}_{\mathbf{I} \cdot \mathbf{I} = \mathbf{I}} \\ &= \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1}\end{aligned}$$

Analysis of **Error** under additive Gaussian noise

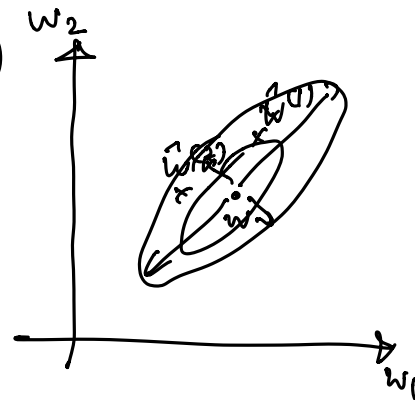
$$\text{if } y_i = x_i^T w + \epsilon_i \quad \text{and} \quad \epsilon_i \stackrel{i.i.d.}{\sim} \mathcal{N}(0, \sigma^2) \quad \mathbf{Y} = \mathbf{X}w + \epsilon$$

$$\begin{aligned} \hat{w}_{MLE} &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y} \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T (\mathbf{X}w + \epsilon) \\ &= w + \underbrace{(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T}_{\text{matrix}} \epsilon \end{aligned}$$

$$\mathbb{E}[\hat{w}_{MLE}] = w$$

$$\text{Cov}(\hat{w}_{MLE}) = \mathbb{E}[(\hat{w} - \mathbb{E}[\hat{w}])(\hat{w} - \mathbb{E}[\hat{w}])^T] = (\mathbf{X}^T \mathbf{X})^{-1} \sigma^2$$

$$\hat{w}_{MLE} \sim \mathcal{N}(w, \underbrace{(\mathbf{X}^T \mathbf{X})^{-1} \sigma^2}_{\text{covariance matrix}})$$



Questions?

① What if $X^T X$ has rank $< d$.
 $d \begin{matrix} \square \\ d \end{matrix}$

$n > d$ & X_i 's in general positions..

$$v^T (X^T X) \cdot v = 0.$$

$$\hat{w}_{MLE}$$

$$(X^T X)^+ \cdot X^T y + c \cdot v$$

$$\hat{w}_{MLE}^{(c)}$$

all in this line active with loss.

