

Lecture 3: Linear regression (continued)



The regression problem in matrix notation

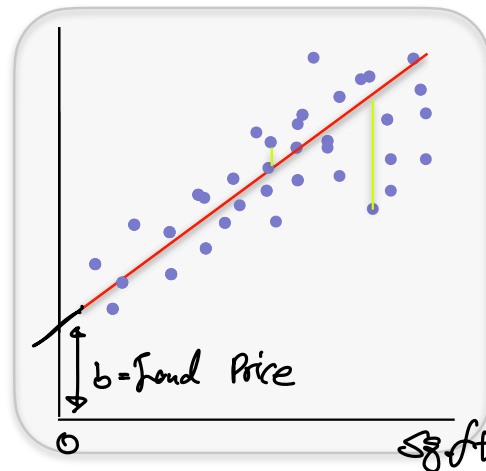
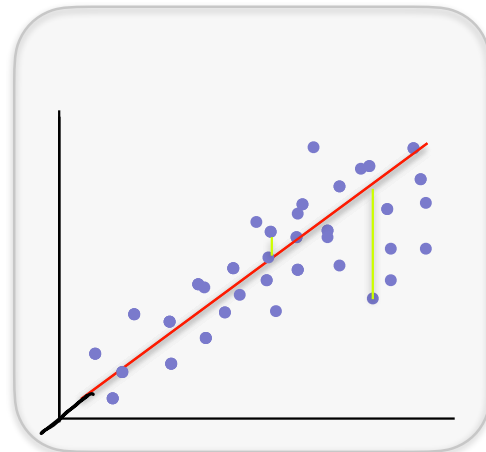
Linear model: $y_i = x_i^T w + \epsilon_i$

$x_i, w \in \mathbb{R}^d$

Least squares solution:

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

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



The regression problem in matrix notation

Linear model: $y_i = x_i^T w + \epsilon_i$
 $x_i, w \in \mathbb{R}^d$

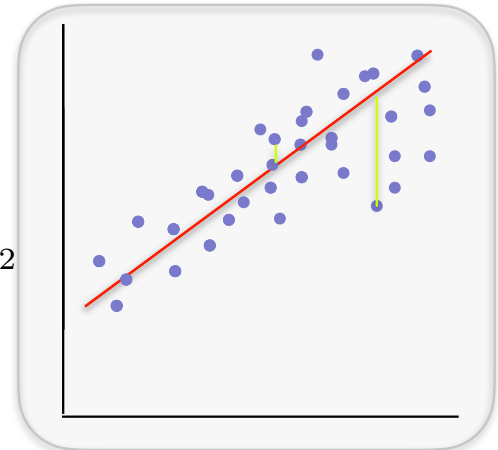
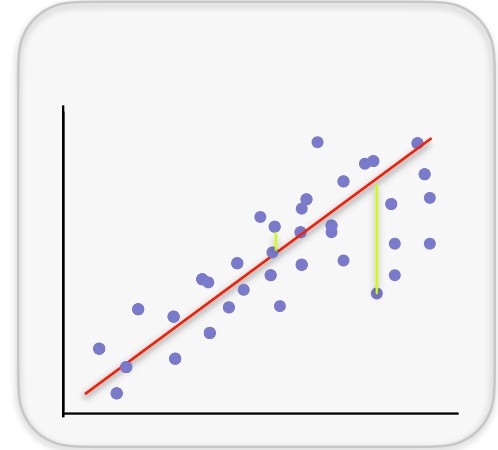
Least squares solution:

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

Affine model: $y_i = x_i^T w + \underbrace{b}_{\substack{\text{Intercept} \\ \mathbb{R}}} + \epsilon_i$

Least squares solution:

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



Dealing with an offset

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

$$=$$

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

$$\begin{aligned} \nabla_w \mathcal{L}(w, b) &\stackrel{?}{=} X^T(y - Xw - \mathbf{1}b) & \in \mathbb{R}^d \\ \nabla_b \mathcal{L}(w, b) &\stackrel{?}{=} \mathbf{1}^T(y - Xw - \mathbf{1}b) & \in \mathbb{R} \end{aligned}$$

$$\begin{aligned} \rightarrow X^T y &= X^T X \hat{w} + X^T \mathbf{1} \hat{b} & \in \mathbb{R}^d \\ \rightarrow \mathbf{1}^T y &= \mathbf{1}^T X \hat{w} + \mathbf{1}^T \mathbf{1} \hat{b} & \in \mathbb{R} \end{aligned}$$

$$\begin{bmatrix} X^T \\ \mathbf{1}^T \end{bmatrix} \cdot y = \begin{bmatrix} X^T \\ \mathbf{1}^T \end{bmatrix} \begin{bmatrix} X & \mathbf{1} \end{bmatrix} \begin{bmatrix} \hat{w} \\ \hat{b} \end{bmatrix}$$

*Series of linear equations with
 $d+1$ -dim variables ($\hat{w}, \hat{b}$)
 $d+1$ equations.

• How do you solve this?

$$\begin{bmatrix} X^T \\ \mathbf{1}^T \end{bmatrix} = \begin{bmatrix} X^T \\ \mathbf{1}^T \end{bmatrix} \begin{bmatrix} X & \mathbf{1} \end{bmatrix} \begin{bmatrix} \hat{w} \\ \hat{b} \end{bmatrix}$$

Schur Complement $\left(\begin{bmatrix} X^T \\ \mathbf{1}^T \end{bmatrix} \begin{bmatrix} X & \mathbf{1} \end{bmatrix} \right)^{-1} \begin{bmatrix} X^T \\ \mathbf{1}^T \end{bmatrix} \cdot y = \begin{bmatrix} \hat{w} \\ \hat{b} \end{bmatrix}$

Dealing with an offset

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

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

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

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

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

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

1th Sample

$$\frac{1}{n} X^T \mathbf{1} = \frac{1}{n} \begin{bmatrix} x_{11} & \dots & x_{1i} & \dots & x_{1n} \\ x_{21} & \dots & x_{2i} & \dots & x_{2n} \\ \vdots & & \vdots & & \vdots \\ x_{n1} & \dots & x_{ni} & \dots & x_{nn} \end{bmatrix} \mathbf{1}$$

1th feature

$$= \begin{bmatrix} \frac{1}{n} \sum_{i=1}^n x_{1i} \\ \vdots \end{bmatrix} = \frac{1}{n} \sum_{i=1}^n x_i = \mu$$

2nd feature

μ

μ_1 : Average Sq. ft.
 μ_2 : Average Bathroom #

Dealing with an offset

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

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

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

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

$$c = b + \mu^T w$$

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

$$\mathcal{L}(w, b) = \|y - Xw + \mathbf{1}\mu^T w - \mathbf{1}\mu^T w - \mathbf{1}b\|_2^2$$

$$\mathcal{L}(w, c) = \|y - (\underbrace{X - \mathbf{1}\mu^T}_{\text{Zero-Mean } \tilde{X}})w - \mathbf{1} \cdot c\|_2^2$$

$$\xrightarrow{\tilde{X}, c}$$

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

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

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

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 $\overline{\mathbf{X}} = [\mathbf{X}, \mathbf{1}]$ and $\overline{w} = \begin{bmatrix} w \\ b \end{bmatrix}$

$$\overline{\mathbf{X}} \overline{w} = \begin{bmatrix} \mathbf{X} & \mathbf{1} \end{bmatrix} \cdot \begin{bmatrix} w \\ b \end{bmatrix} = \mathbf{X} \cdot w + \mathbf{1} \cdot b$$

$$\begin{array}{c} \begin{array}{cc} & d+1 \\ \boxed{\begin{array}{c} \mathbf{X} \quad \mathbf{1} \end{array}} & \boxed{\begin{array}{c} w \\ b \end{array}} \\ n \end{array}$$

$$\hat{\overline{w}}_{LS} = \arg \min \| \mathbf{y} - \overline{\mathbf{X}} \cdot \overline{w} \|_2^2 = (\overline{\mathbf{X}}^T \overline{\mathbf{X}})^{-1} \cdot \overline{\mathbf{X}}^T \cdot \mathbf{y}$$

$$\begin{bmatrix} \hat{w} \\ \hat{b} \end{bmatrix} = \left(\begin{bmatrix} \mathbf{X}^T \\ \mathbf{1}^T \end{bmatrix} \begin{bmatrix} \mathbf{X} & \mathbf{1} \end{bmatrix} \right)^{-1} \cdot \begin{bmatrix} \mathbf{X} & \mathbf{1} \end{bmatrix} \cdot \mathbf{y}$$

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 $\epsilon_i \stackrel{i.i.d.}{\sim} \mathcal{N}(0, \sigma^2)$

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

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

Why is least squares a good loss function?

Maximum Likelihood Estimator:

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

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

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

* Gaussian Noise $\iff$ ℓ_2 -loss minimization.

Recap of linear regression

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

Minimize the loss (Empirical Risk Minimization)

Choose a loss
e.g., $(y_i - x_i^T w)^2$ *there are other choices*

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

Maximize the likelihood (MLE)

there are other hypotheses.
Choose a Hypothesis class
e.g., $y_i = x_i^T w + \epsilon_i, \epsilon_i \sim \mathcal{N}(0, \sigma^2)$

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

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

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

$\uparrow$ Random

Maximum Likelihood Estimator is unbiased:

$$\begin{aligned}\mathbb{E}[\hat{w}_{MLE}] &= w^* + \mathbb{E}[(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \epsilon] \quad \leftarrow \text{Linearity of expectation} \\ &= w^* + (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \cdot \underbrace{\mathbb{E}[\epsilon]}_0 \quad \leftarrow \mathbb{E}[a \cdot \epsilon + b] = a \mathbb{E}[\epsilon] + b \\ &= w^*\end{aligned}$$

Analysis of **Error** under additive Gaussian noise

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

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

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

$$= \mathbb{E}[(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \epsilon \epsilon^T \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1}]$$

$$= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbb{E}[\epsilon \epsilon^T] \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1}$$

$$= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \sigma^2 \mathbf{I} \cdot \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1}$$

$$= \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1}$$

$$= \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1}$$

linearity of
expectation

$$\begin{aligned}\mathbb{E}[\epsilon \epsilon^T] &= \mathcal{M} \\ \mathcal{M}_{ij} &= \mathbb{E}[\epsilon_i \epsilon_j] \\ &= \begin{cases} \sigma^2 & \text{if } i=j \\ 0 & \text{else} \end{cases} \\ \mathcal{M} &= \begin{bmatrix} \sigma^2 & 0 & \dots & 0 \\ 0 & \sigma^2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \sigma^2 \end{bmatrix} = \sigma^2 \mathbf{I}.\end{aligned}$$

Analysis of **Error** under additive Gaussian noise

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

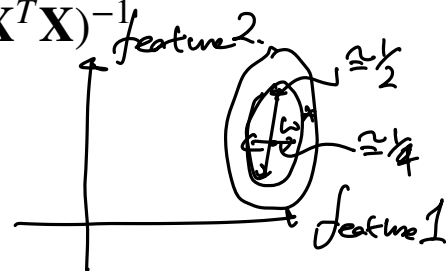
$$\begin{aligned}\hat{w}_{MLE} &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y} \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T (\mathbf{X}w + \epsilon) \\ &= w + (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \epsilon\end{aligned}$$

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

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

$$\hat{w}_{MLE} \sim \mathcal{N}(w, \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1}) = \mathcal{N}(w^*, \sigma^2 \begin{bmatrix} \frac{1}{4} & 0 \\ 0 & \frac{1}{2} \end{bmatrix})$$

If $X = \begin{bmatrix} 1 & 0 \\ 1 & 0 \\ 1 & 0 \\ 1 & 0 \\ 0 & 1 \end{bmatrix}$ then $X^T X = \begin{bmatrix} 4 & 0 \\ 0 & 2 \end{bmatrix}$ $(X^T X)^{-1} = \begin{bmatrix} \frac{1}{4} & 0 \\ 0 & \frac{1}{2} \end{bmatrix}$



Questions?
