

Maximum Likelihood Estimation

Observe $X_1, X_2, \dots, X_n$ drawn IID from $f(x; \theta)$ for some “true” $\theta = \theta_*$

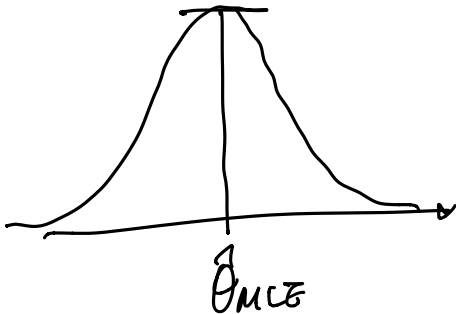
$\theta \in \mathcal{H}$
 $1-\theta \in \mathcal{T}$

Likelihood function $L_n(\theta) = \prod_{i=1}^n f(X_i; \theta) = \theta^K (1-\theta)^{n-K}$

Log-Likelihood function $l_n(\theta) = \log(L_n(\theta)) = \sum_{i=1}^n \log(f(X_i; \theta)) = K \log \theta + (n-K) \log(1-\theta)$

Maximum Likelihood Estimator (MLE) $\hat{\theta}_{MLE} = \arg \max_{\theta} L_n(\theta)$

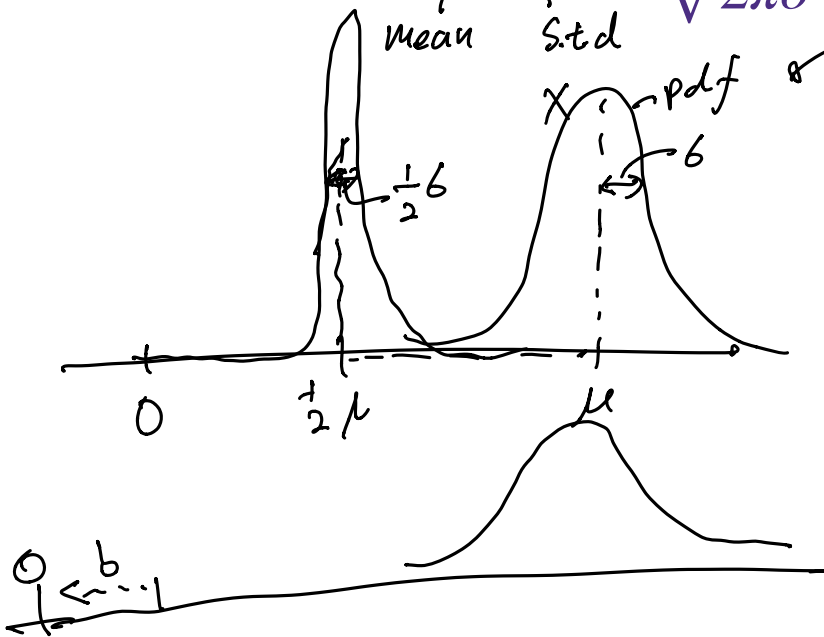
$= \frac{K}{n}$



What about continuous variables?

- *Client*: What if I am measuring a **continuous variable**?
- *You*: Let me tell you about Gaussians...

$$P(x; \mu, \sigma) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$



$$X \sim N(\mu, b^2)$$

$$X+b \quad \left| \quad \frac{1}{2}X\right.$$

Some properties of Gaussians

- affine transformation (multiplying by scalar and adding a constant)
 - $X \sim N(\mu, \sigma^2)$
 - $Y = aX + b \rightarrow Y \sim N(a\mu + b, a^2\sigma^2)$
- Sum of Gaussians
 - $X \sim N(\mu_X, \sigma_X^2)$
 - $Y \sim N(\mu_Y, \sigma_Y^2)$
 - $Z = X + Y \rightarrow Z \sim N(\mu_X + \mu_Y, \sigma_X^2 + \sigma_Y^2)$

MLE for Gaussian

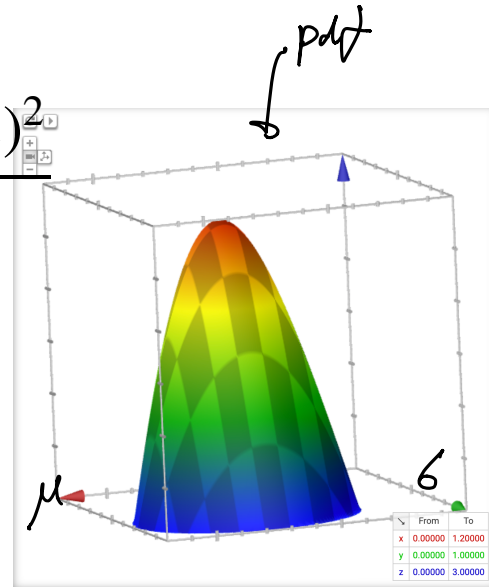
- Prob. of i.i.d. samples $D=\{x_1, \dots, x_n\}$ (e.g., temperature):

$$\begin{aligned}
 P(\mathcal{D}; \mu, \sigma) &= P(x_1, \dots, x_n; \mu, \sigma) = p(x_1 | \mu, \sigma) \times \dots \times p(x_n | \mu, \sigma) \\
 &= \prod_{i=1}^n \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{(x_i - \mu)^2}{2\sigma^2}} = \underbrace{\prod_{i=1}^n p(x_i | \mu, \sigma)}_{\text{pdf}}
 \end{aligned}$$

- Log-likelihood of data:

$$\log P(\mathcal{D}; \mu, \sigma) = -n \log(\sigma \sqrt{2\pi}) - \sum_{i=1}^n \frac{(x_i - \mu)^2}{2\sigma^2}$$

- What is $\hat{\theta}_{MLE}$ for $\theta = (\mu, \sigma^2)$?



Your second learning algorithm: MLE for mean of a Gaussian

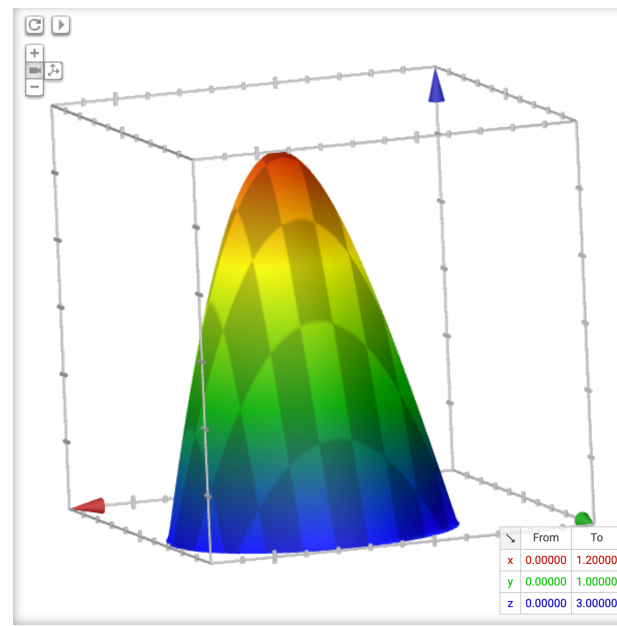
- What's MLE for mean?

$$\frac{d}{d\mu} \log P(\mathcal{D}; \mu, \sigma) = \frac{d}{d\mu} \left[-n \log(\sigma\sqrt{2\pi}) - \sum_{i=1}^n \frac{(x_i - \mu)^2}{2\sigma^2} \right]$$

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

$$= \frac{1}{\sigma^2} \sum_{i=1}^n (x_i - \mu) = \frac{1}{\sigma^2} \left(\sum_{i=1}^n x_i - n\mu \right) \Big|_{\mu = \hat{\mu}} = 0$$

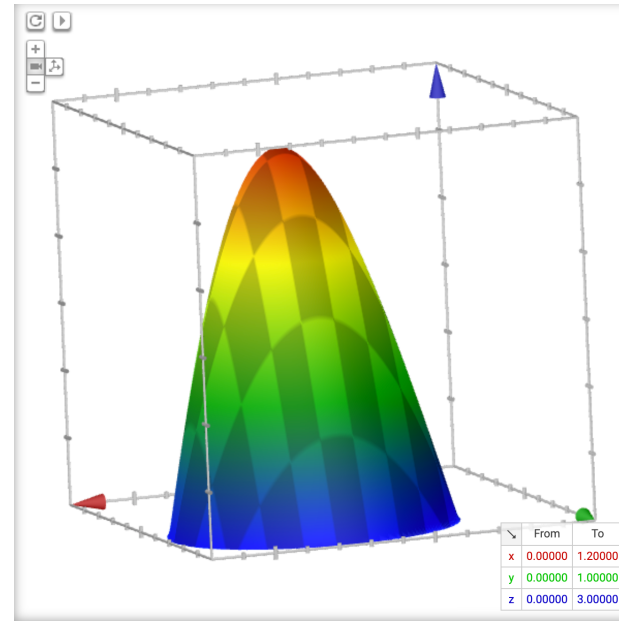
$$\frac{1}{n} \sum_{i=1}^n x_i = \hat{\mu}_{LS}$$



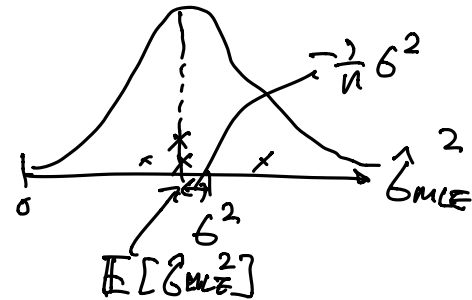
MLE for variance

- Again, set derivative to zero:

$$\begin{aligned}\frac{d}{d\sigma} \log P(\mathcal{D}; \mu, \sigma) &= \frac{d}{d\sigma} \left[-n \log(\sigma\sqrt{2\pi}) - \sum_{i=1}^n \frac{(x_i - \mu)^2}{2\sigma^2} \right] \\&= \frac{-n}{\sigma} - \sum_{i=1}^n \frac{(x_i - \mu)^2 (-2)}{2\sigma^3} \\&= \frac{1}{\sigma^3} \left(-n\sigma^2 + \sum_{i=1}^n (x_i - \mu)^2 \right) \bigg|_{\substack{\mu = \hat{\mu} \\ \sigma = \hat{\sigma}}} = 0 \\-n\hat{\sigma}^2 + \sum_{i=1}^n (x_i - \hat{\mu})^2 &= 0 \\ \hat{\sigma}^2 &= \frac{1}{n} \sum_{i=1}^n (x_i - \hat{\mu})^2\end{aligned}$$



What can we say about the MLE?



- MLE:

mean
 μ

variance
 σ^2

$$\hat{\mu}_{MLE} = \frac{1}{n} \sum_{i=1}^n x_i$$

$$, \mathbb{E}_{x_i \sim N(\mu, \sigma)} [\hat{\mu}_{MLE}] = \mu$$

$\mathbb{E}$	$\checkmark$	$\checkmark$
VAR	?	?

$$\hat{\sigma}_{MLE}^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \hat{\mu}_{MLE})^2$$

- MLE for the variance of a Gaussian is **biased**

$$\frac{n-1}{n} \cdot \sigma^2 = \mathbb{E}[\hat{\sigma}_{MLE}^2] \neq \sigma^2$$

$$(1 - \frac{1}{n}) \sigma^2 = \mathbb{E}[\hat{\sigma}_{MLE}^2]$$

$$\text{bias} = -\frac{1}{n} \sigma^2$$

- Unbiased variance estimator:

$$\hat{\sigma}_{unbiased}^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \hat{\mu}_{MLE})^2$$

Maximum Likelihood Estimation

$$\frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x-\mu)^2}{2\sigma^2}} \quad \theta = (\mu, \sigma)$$

Observe $X_1, X_2, \dots, X_n$ drawn IID from $f(x; \theta)$ for some “true” $\theta = \theta_*$

Likelihood function $L_n(\theta) = \prod_{i=1}^n f(X_i; \theta)$

Log-Likelihood function $l_n(\theta) = \log(L_n(\theta)) = \sum_{i=1}^n \log(f(X_i; \theta))$

Maximum Likelihood Estimator (MLE) $\hat{\theta}_{MLE} = \arg \max_{\theta} L_n(\theta)$

Properties (under benign regularity conditions—smoothness, identifiability, etc.):

- Asymptotically consistent and normal: $\frac{\hat{\theta}_{MLE} - \theta_*}{\widehat{se}} \sim \mathcal{N}(0, 1)$
- Asymptotic Optimality, minimum variance (see Cramer-Rao lower bound)

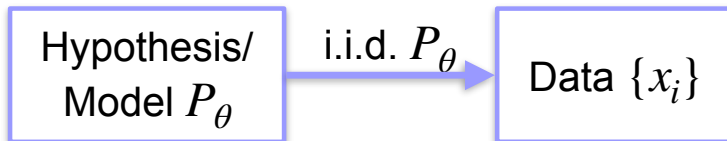
Recap

- Learning is...
 - Collect some data
 - E.g., coin flips

Data $\{x_i\}$

Recap

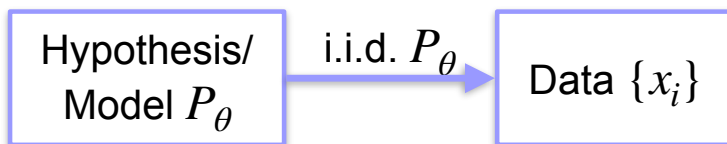
- Learning is...
 - Collect some data
 - E.g., coin flips
 - Choose a hypothesis class or model
 - E.g., binomial



Recap

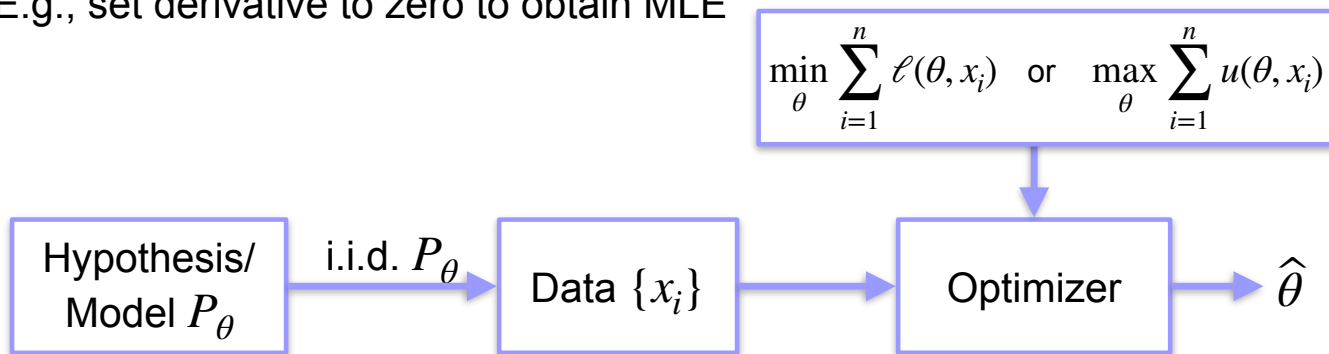
- Learning is...
 - Collect some data
 - E.g., coin flips
 - Choose a hypothesis class or model
 - E.g., binomial
 - Choose a loss function
 - E.g., data likelihood

$$\min_{\theta} \sum_{i=1}^n \ell(\theta, x_i) \quad \text{or} \quad \max_{\theta} \sum_{i=1}^n u(\theta, x_i)$$



Recap

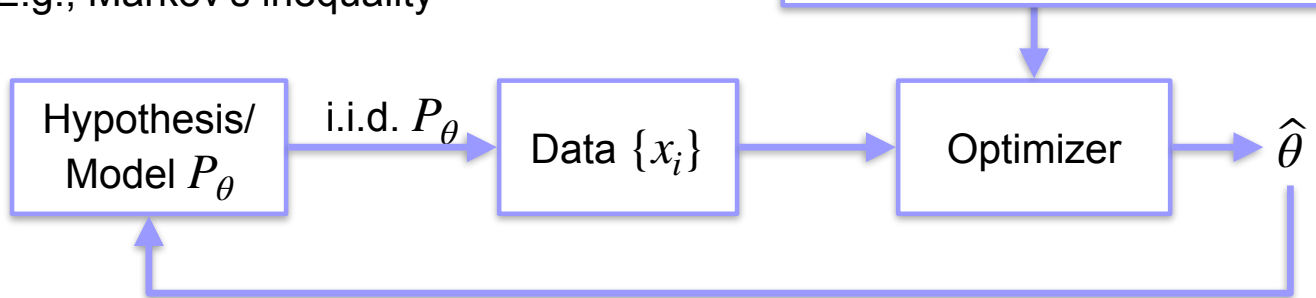
- Learning is...
 - Collect some data
 - E.g., coin flips
 - Choose a hypothesis class or model
 - E.g., binomial
 - Choose a loss function
 - E.g., data likelihood
 - Choose an optimization procedure
 - E.g., set derivative to zero to obtain MLE



Recap

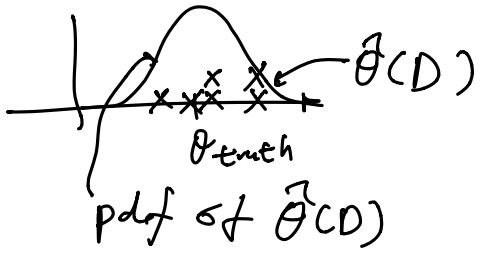
- Learning is...
 - Collect some data
 - E.g., coin flips
 - Choose a hypothesis class or model
 - E.g., binomial
 - Choose a loss function
 - E.g., data likelihood
 - Choose an optimization procedure
 - E.g., set derivative to zero to obtain MLE
 - Justifying the accuracy of the estimate
 - E.g., Markov's inequality

$$\min_{\theta} \sum_{i=1}^n \ell(\theta, x_i) \quad \text{or} \quad \max_{\theta} \sum_{i=1}^n u(\theta, x_i)$$



usually, $D = \{x_i\}_{i=1}^n \rightarrow \frac{\hat{\theta}(D)}{\text{Random variable}}$

Imagine, again,



pdf of $\hat{\theta}(D)$

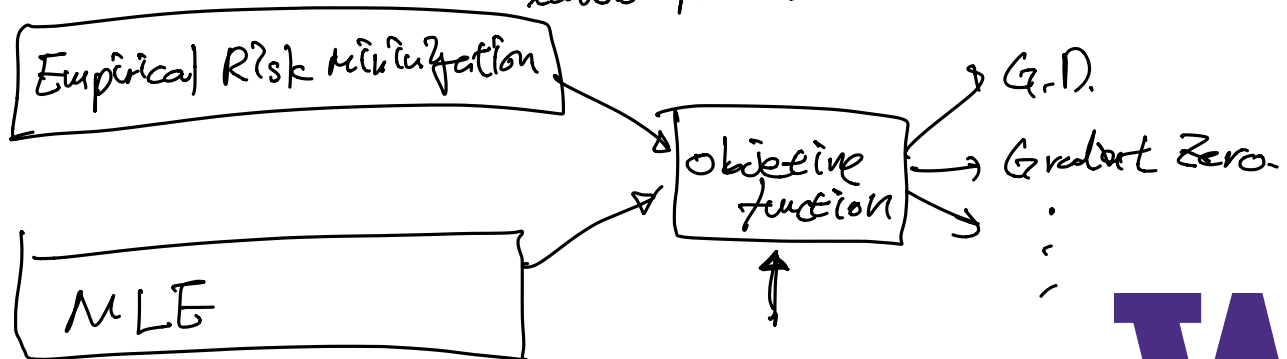
Hypothesis class is linear function

$$\sum_{i=1}^d w_i \cdot x_{1,i} = y_i$$

||

Linear Regression

||
labels predicted real-valued



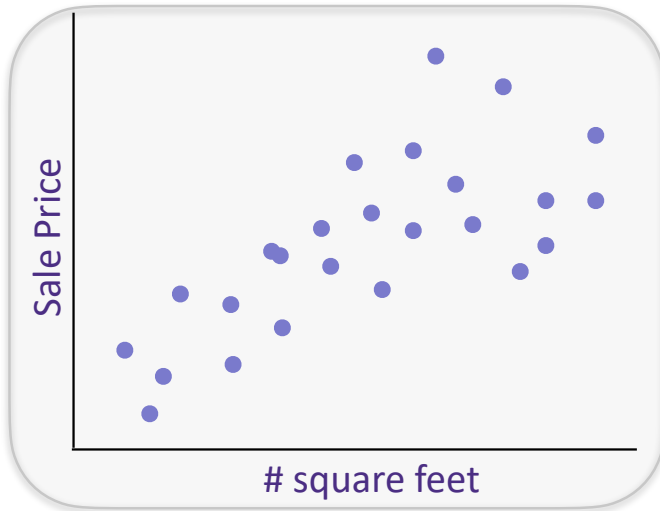
W

The regression problem, 1-dimensional

Given past sales data on [zillow.com](https://www.zillow.com), predict:

y = House sale price from

x = {# sq. ft.}



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

$x_i \in \mathbb{R}$
 $y_i \in \mathbb{R}$

Process

Decide on a **model**

assume y_i house sale price is a linear function of
square feet.
 x_i

Find the function which fits the data best

Use function to make prediction on new examples

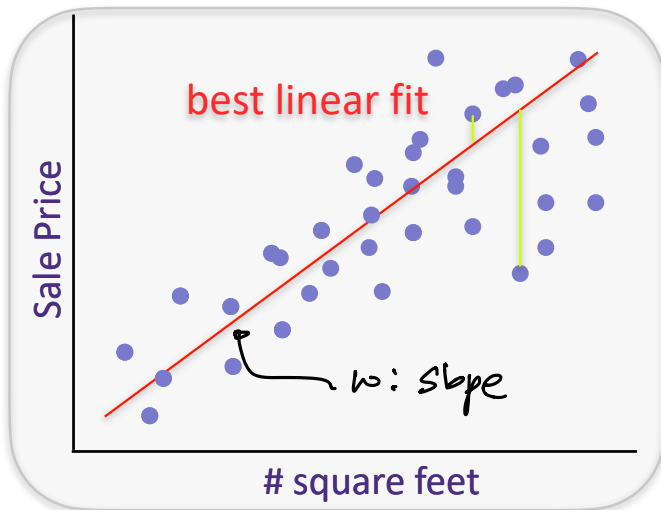
$$y_{\text{new}} = f(x_{\text{new}})$$

Fit a function to our data, 1-dimension

Given past sales data on [zillow.com](https://www.zillow.com), predict:

y = House sale price *from*

x = {# sq. ft.}



Error

$$y_i = x_i w + \epsilon_i$$

Training Data: $x_i \in \mathbb{R}$
 $y_i \in \mathbb{R}$
 $\{(x_i, y_i)\}_{i=1}^n$

Hypothesis/Model: linear

$$y_i \approx x_i w$$

Loss: least squares solution

$$\min_w \sum_{i=1}^n \underbrace{(y_i - x_i w)^2}_{\epsilon_i: \text{error}}$$

The regression problem, d-dimensions

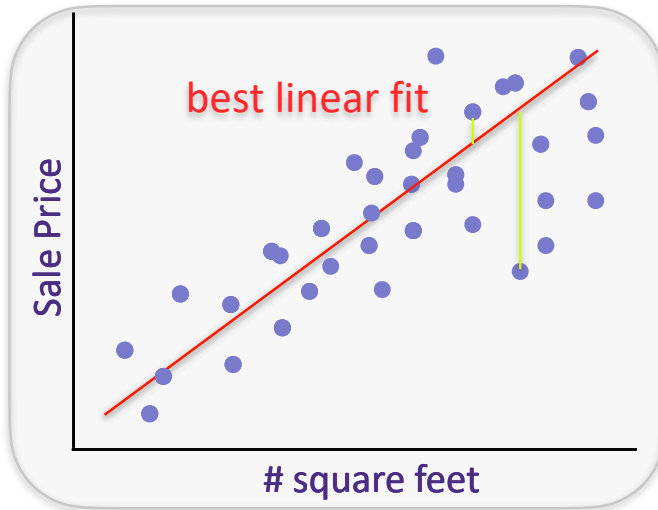
Given past sales data on [zillow.com](https://www.zillow.com), predict:

y = House sale price *from*

x = {# sq. ft., zip code, date of sale, etc.}

Error:

$$y_i = x_i^T w + \epsilon_i$$



Training Data: $x_i \in \mathbb{R}^d$
 $\{(x_i, y_i)\}_{i=1}^n$ $y_i \in \mathbb{R}$
 $w \in \mathbb{R}^d$

Hypothesis/Model: linear

$$y_i \approx x_i^T w = x_{i,1} w_1 + x_{i,2} w_2 + \dots + x_{i,d} w_d$$

Loss: least squares solution

$$\min_{\substack{w \\ w \in \mathbb{R}^d}} \sum_{i=1}^n \underbrace{(y_i - x_i^T w)}_{\text{error}}^2$$

The regression problem in matrix notation

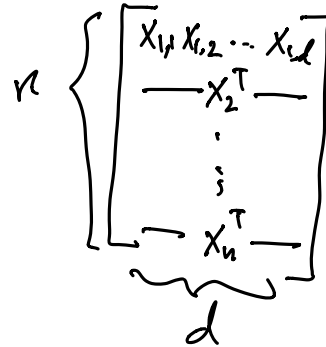
Data:

$$\mathbf{y} = \left[\begin{array}{c} y_1 \\ \vdots \\ y_n \end{array} \right] \quad \mathbf{X} = \left[\begin{array}{c} x_1^T \\ \vdots \\ x_n^T \end{array} \right]$$

d : # of features

n : # of examples/datapoints

$$n > d$$



A hand-drawn diagram of the matrix $\mathbf{X}$. It is a large square bracket containing a grid of elements. The top row is labeled $x_{1,1}, x_{1,2}, \dots, x_{1,d}$. Below this, there is a row labeled x_2^T , followed by a vertical ellipsis, and then a row labeled x_n^T . A large curly brace on the left side of the matrix is labeled n . A curly brace at the bottom of the matrix is labeled d .

The regression problem in matrix notation

Data:

$$\mathbf{y} = \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} \quad \mathbf{X} = \begin{bmatrix} x_1^T \\ \vdots \\ x_n^T \end{bmatrix}$$

d : # of features

n : # of examples/datapoints

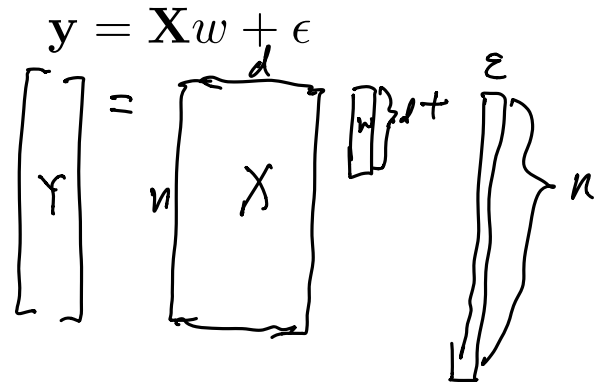
Model:

$$y_1 = x_1^T w + \epsilon_1$$

$$y_2 = x_2^T w + \epsilon_2$$

$\vdots$

$$y_n = x_n^T w + \epsilon_n$$

$$\mathbf{y} = \mathbf{X}w + \epsilon$$


The regression problem in matrix notation

Data:

$$\mathbf{y} = \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} \quad \mathbf{X} = \begin{bmatrix} x_1^T \\ \vdots \\ x_n^T \end{bmatrix}$$

d : # of features

n : # of examples/datapoints

Model:

$$y_1 = x_1^T w + \epsilon_1$$

$$y_2 = x_2^T w + \epsilon_2$$

$\vdots$

$$y_n = x_n^T w + \epsilon_n$$

$$\mathbf{y} = \mathbf{X}w + \epsilon$$

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

$\leftarrow \sum_i v_i^2 = \|v\|_2^2 = V^T \cdot V$

$\|v\|_2 = \sqrt{v_1^2 + v_2^2 + \dots}$

The regression problem in matrix notation

Data:

$$\mathbf{y} = \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} \quad \mathbf{X} = \begin{bmatrix} x_1^T \\ \vdots \\ x_n^T \end{bmatrix}$$

d : # of features

n : # of examples/datapoints

Model:

$$y_1 = x_1^T w + \epsilon_1$$

$$y_2 = x_2^T w + \epsilon_2$$

$\vdots$

$$y_n = x_n^T w + \epsilon_n$$

$$\mathbf{y} = \mathbf{X}w + \epsilon$$

$$\begin{bmatrix} y_1 - x_1^T w \\ \vdots \\ y_n - x_n^T w \end{bmatrix} = \mathbf{y} - \mathbf{X}w$$

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

$$= \arg \min_w \|\mathbf{y} - \mathbf{X}w\|_2^2$$

$$= \arg \min_w (\mathbf{y} - \mathbf{X}w)^T (\mathbf{y} - \mathbf{X}w)$$

$$\begin{aligned} \|\mathbf{v}\|_2 &\triangleq \sqrt{v_1^2 + \dots + v_d^2} \\ \|\mathbf{v}\|_2^2 &= (\|\mathbf{v}\|_2)^2 \\ &= v_1^2 + \dots + v_d^2 \\ &= \mathbf{v}^T \mathbf{v} \end{aligned}$$

The regression problem in matrix notation

parameter $\theta = w$

$$\hat{w}_{LS} = \arg \min_{w \in \mathbb{R}^d} \|y - Xw\|_2^2$$

$$= \arg \min_w (y - Xw)^T (y - Xw)$$

$\mathcal{L}(w)$

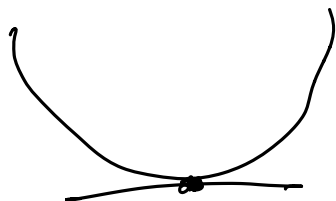
Set gradient w.r.t. w to zero to find the minima:

$$\frac{1}{2} \nabla_w \mathcal{L}(w) = -X^T (y - Xw)$$

$$= -X^T y + X^T X w \Big|_{w = \hat{w}_{LS}} = 0$$

$$X^T X \cdot \hat{w} = X^T y$$

$$\underbrace{d \times d}_{\text{invertible}} \cdot d = d \times n \cdot n$$



$$f(w) = w^T w = w_1^2 + w_2^2 + \dots$$

$$\nabla_w f(w) = 2 \cdot w$$

$$f(w) = (Aw)^T (Aw)$$

$$\nabla_w f(w) = 2A^T A w$$

$$f(w) = (Aw + b)^T (Aw + b)$$

$$\nabla_w f(w) = 2A^T (Aw + b)$$

$$\|v\|_p = (|v_1|^p + \dots)^{1/p}$$

$$p = 0, 1, 2, \infty$$

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

$$\square \cdot \square$$

closed-form

The regression problem in matrix notation

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

“Closed form” solution!