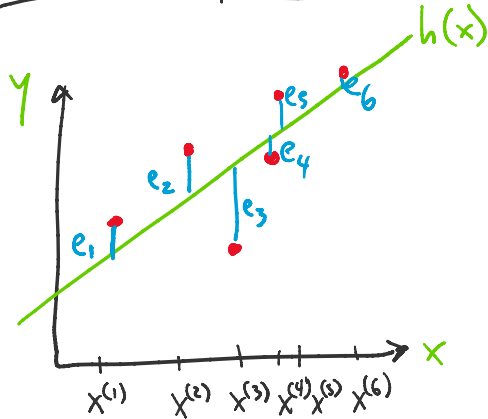


Section B: Linear Regression + Gradient Descent

Wednesday, September 21, 2022 12:49 PM

Residuals $D = \{(\vec{x}^{(i)}, y^{(i)})\}_{i=1}^N$ $\hat{y} = h_{\vec{\theta}}(\vec{x})$

Def: a residual is the vertical distance from observed output $y^{(i)}$ to predicted output $\hat{y} = h(\vec{x}^{(i)})$



$$e_i = |y^{(i)} - h(\vec{x}^{(i)})|$$

$$= |y^{(i)} - (\vec{w}^T \vec{x}^{(i)} + b)|$$

Key Idea for Lin. Reg.

find the linear function h (w/parameters $\vec{w}$ and b) that minimizes the squares of the residuals for a training dataset

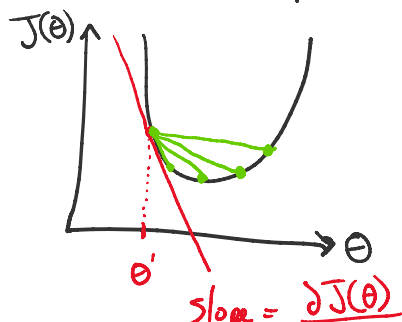
one option among many

Def: mean squared error (MSE)

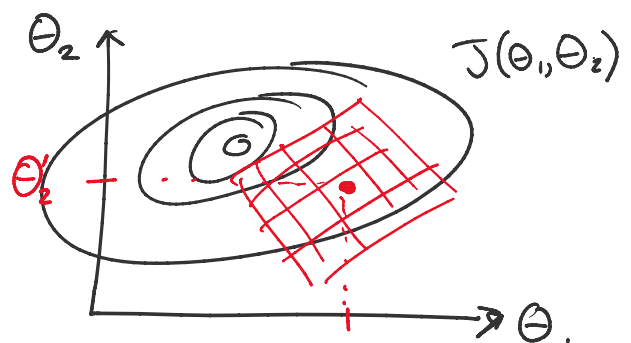
$$J_D(\vec{w}, b) = \frac{1}{N} \sum_{i=1}^N (e_i)^2 = \frac{1}{N} \sum_{i=1}^N (y^{(i)} - (\vec{w}^T \vec{x}^{(i)} + b))^2$$

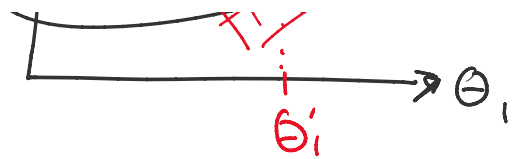
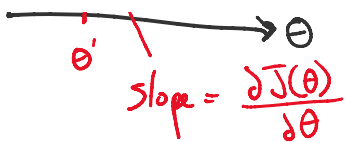
Derivatives

① Deriv. as Slope of Tangent



③ Partial Deriv. as Tangent Plane





(2) Deriv. as a Limit

$$\frac{\partial J(\theta)}{\partial \theta} = \lim_{e \rightarrow 0} \frac{J(\theta+e) - J(\theta)}{e}$$

limit of secants is the tangent

$$\begin{bmatrix} \frac{\partial J(\theta_1, \theta_2)}{\partial \theta_1} \\ \frac{\partial J(\theta_1, \theta_2)}{\partial \theta_2} \end{bmatrix}$$

Gradient

Def: the gradient of $J: \mathbb{R}^M \rightarrow \mathbb{R}$

$$\nabla J(\vec{\theta}) = \begin{bmatrix} \partial J(\vec{\theta}) / \partial \theta_1 \\ \partial J(\vec{\theta}) / \partial \theta_2 \\ \vdots \\ \partial J(\vec{\theta}) / \partial \theta_M \end{bmatrix}$$

first order partial derivative

Gradient Descent

Algorithm:

(1) Choose an initial $\vec{\theta}$

(2) Repeat $t=1, 2, 3, \dots$

a) Compute gradient $\vec{g} = \nabla J(\vec{\theta})$

b) Choose a step size γ_t

c) Update $\vec{\theta} \leftarrow \vec{\theta} - \gamma_t \vec{g}$

(3) Return $\vec{\theta}$ when stopping criterion is reached

Initial Point

Step Size

Stopping Criterion

Initial Point

- randomly
- $\vec{\theta} = \text{all zeros}$

Step Size

- fixed value $\gamma = 0.1$
- schedule $\gamma_t = \frac{\gamma_0}{(t-1)\gamma + 1}$

Stopping Criterion

- when $\vec{g} \approx 0$
- $\|\nabla J(\theta)\|_2 < \epsilon$
for $\epsilon = 10^{-8}$

Gradient for Lin. Reg.

MSE: $J(\vec{\theta}) = \frac{1}{N} \sum_{i=1}^N J^{(i)}(\theta)$ where $J^{(i)}(\theta) = \frac{1}{2} (y^{(i)} - (\vec{\theta}^T \vec{x}^{(i)}))^2$

does not change the argmin

Partial Deriv.s

$$\begin{aligned} \frac{\partial J^{(i)}(\vec{\theta})}{\partial \theta_j} &= \cancel{\frac{1}{2}} \cdot 2 (y^{(i)} - \vec{\theta}^T \vec{x}^{(i)}) \frac{\partial}{\partial \theta_j} (y^{(i)} - \vec{\theta}^T \vec{x}^{(i)}) \\ &= (y^{(i)} - \vec{\theta}^T \vec{x}^{(i)}) \frac{\partial}{\partial \theta_j} (y^{(i)} - \sum_{m=1}^M \theta_m x_m^{(i)}) \\ &= - \underbrace{(y^{(i)} - \vec{\theta}^T \vec{x}^{(i)})}_{\text{scalar}} x_j^{(i)} \end{aligned}$$

Gradient:

$$\nabla J^{(i)}(\vec{\theta}) = \begin{bmatrix} \partial J^{(i)}(\theta) / \partial \theta_1 \\ \vdots \\ \partial J^{(i)}(\theta) / \partial \theta_M \end{bmatrix} = \underbrace{(y^{(i)} - \vec{\theta}^T \vec{x}^{(i)})}_{\text{scalar}} \underbrace{\vec{x}^{(i)}}_{\text{vector}}$$

$$\nabla J(\theta) = \nabla \left(\frac{1}{N} \sum_{i=1}^N J^{(i)}(\theta) \right) = \frac{1}{N} \sum_{i=1}^N \nabla J^{(i)}(\theta)$$

$$= \frac{1}{N} \sum_{i=1}^N - (y^{(i)} - \vec{\theta}^T \vec{x}^{(i)}) \vec{x}^{(i)}$$

known b/c
dataset

known b/c
we passed it in

$$D = \{(\vec{x}^{(1)}, y^{(1)}), (\vec{x}^{(2)}, y^{(2)}), \dots, (\vec{x}^{(N)}, y^{(N)})\}$$

$$\vec{x}^{(2)} = [x_1^{(2)} \ x_2^{(2)} \ \dots \ x_M^{(2)}]$$

$$h(\vec{x}) = \vec{\theta}^T \vec{x}$$

$$y = w_1 x_1 + b$$

$$y = w_1 x_1 + w_2 x_2 + b$$

$$y = \vec{w}^T \vec{x} + b$$

$$y = \underbrace{\theta_1}_{\substack{\text{intercept} \\ = 1}} x_1 + \theta_2 x_2$$

$$y = \theta_1 x_1 + \theta_2 x_2 + \theta_3 x_3$$

$$y = \vec{\theta}^T \vec{x}$$

slope
 intercept
 slope₂
 slope₃