

$\times$ $X_{\text{-train}}: m \times n$

$$\begin{matrix}
 & n \\
 m & \left[\begin{matrix} & & & \\ & & & \\ & & & \end{matrix} \right]
 \end{matrix}
 \rightarrow \text{row} = \text{a training data point}$$
 There're m training data points
 $\uparrow$ column = a feature

$\vec{w}$ weights: $n \times 1$ a parameter for each feature

$\vec{y}$ $y_{\text{-train}}: m \times 1$ a output / label for each training data pt.

$\vec{e}$ error: $m \times 1$ same size as $y_{\text{-train}}$

$\hat{\vec{y}}$ output = $X \vec{w} + \vec{b}$ $\vec{b}: m \times 1$ bias for each prediction

$\vec{e} = \frac{1}{m} (\hat{\vec{y}} - \vec{y})^2$
 mean square error, vector square is $y^T y$

$= \frac{1}{m} (\hat{\vec{y}} - \vec{y})^T (\hat{\vec{y}} - \vec{y})$
 $\hat{\vec{y}}^T \vec{y} = \vec{y}^T \hat{\vec{y}}$ b/c outcome is scalar

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

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

$\nabla_{\vec{w}} \vec{e} = \begin{matrix} \downarrow & \downarrow & \downarrow & \downarrow & \downarrow & \downarrow \\ 2X^T X \vec{w} & X^T \vec{b} & X^T \vec{b} & 0 & -2X^T \vec{y} & 0 \end{matrix}$

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

$= \frac{2}{m} X^T \vec{e} \quad (n \times 1)$

$\nabla_{\vec{b}} \vec{e} = \frac{1}{m} (X \vec{w} + X \vec{w} + 2\vec{b}) = \frac{2}{m} (X \vec{w} + \vec{b})$

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

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

$\nabla_{\vec{w}} (\vec{w}^T A \vec{w}) = (A + A^T) \vec{w}$

$$\hat{\vec{y}} = \vec{x}' \vec{w} \quad \vec{e} = \vec{x}' \vec{w} - \vec{y}$$

$$\vec{e} = \frac{1}{m} (\vec{w}^T \vec{x}' \vec{x}'^T \vec{w} - 2 \vec{y}^T \vec{x}' \vec{w})$$

$\vec{x}' = m \times (n+1)$ last col $\begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix}$
 $\vec{w} = (n+1) \times 1$
the $\vec{b}$ is just $\vec{w}$

$$\nabla_{\vec{w}} \vec{e} = \frac{1}{m} (2 \vec{x}'^T \vec{x}' \vec{w} - 2 \vec{x}'^T \vec{y})$$

$$= \frac{2}{m} (2 \vec{x}'^T (\vec{x}' \vec{w} - \vec{y}))$$

$$= \frac{2}{m} \vec{x}'^T \vec{e} \quad n+1 \times 1$$

$$\left[\vec{x} \mid \vec{b} \right] \begin{bmatrix} \vec{w} \end{bmatrix}$$

Takeaway:

- weights correspond to the dimensions in training data, bias is the same across all training data, it's the y-intercept

- stack a column of 1s at the end of matrix X to easily calculate $\vec{x}' \vec{w} + \vec{b}$ with $\vec{x}' \vec{w}$

$$m \begin{bmatrix} \vec{x} \mid \begin{matrix} \uparrow \\ \vdots \\ \uparrow \end{matrix} \end{bmatrix} \begin{matrix} \rightarrow \text{each training data} \\ \text{bias is the same across all training data,} \\ \text{each dimension / feature.} \end{matrix}$$

- $\vec{w}$ has length = $n+1$. $\nabla_{\vec{w}}$ finds gradient for each w_i and gradient for b . w_i stands for slope for each dimension, b is y intercept

- vector calculus:

$$\left. \begin{aligned} \nabla_{\vec{w}} (\vec{x}' \vec{w}) &= \vec{x} \\ \nabla_{\vec{w}} (\vec{w}^T \vec{x}) &= \vec{x} \end{aligned} \right\} \text{ b/c } \vec{x}' \vec{w} \text{ is a scalar, } \vec{x}' \vec{w} = \vec{w}^T \vec{x}$$

$$\nabla_{\vec{w}} (\vec{w}^T A \vec{w}) = (A + A^T) \vec{w}$$

- Use mean square error for loss, b/c it gives us a convex function that we can optimize over by gradient descent
- Math for Linear Regression w Gradient Descent

X_{train} = matrix of training data $m \times n$

x = X_{train} stack w/ a column of 1s $m \times (n+1)$

$\vec{w}$ = length $n+1$

$$\vec{\hat{y}} = x \vec{w} \quad \vec{e} = \vec{y} - \vec{\hat{y}} \quad \text{error} = \text{actual} - \text{predicted}$$

$$\begin{aligned} l &= \frac{1}{m} (\vec{y} - \vec{\hat{y}})^2 = \frac{1}{m} (\vec{y} - x \vec{w})^2 = \frac{1}{m} (\vec{y} - x \vec{w})^T (\vec{y} - x \vec{w}) \\ &= \frac{1}{m} (\vec{w}^T x^T x \vec{w} - 2 \vec{w}^T x^T \vec{y} + \vec{y}^T \vec{y}) \end{aligned}$$

$$\nabla_{\vec{w}} l = \frac{1}{m} [(x^T x + x x^T) \vec{w} - 2 x^T \vec{y} + 0]$$

$$= \frac{1}{m} (2 x^T x \vec{w} - 2 x^T \vec{y})$$

$$= \frac{2}{m} x^T (x \vec{w} - \vec{y}) = -\frac{2}{m} x^T \vec{e}$$

$\nabla_{\vec{w}} l$ is used for gradient descent! step opposite of gradient direction

- Recall Linear Regression Normal form Equation

$$x = \left[\boxed{\text{data}} \middle| \begin{matrix} 1 \\ \vdots \\ 1 \end{matrix} \right] \quad \vec{y} = \text{labels}$$

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

High Level Notes:

- Regression Problems: y are continuous values $\rightarrow$ Lin Reg
Classification problems: y are binary $\rightarrow$ Logistic Regression
- Lasso & Ridge Regularization are applied to linear Regression so that the weights are minimized.

Regularization prevents overfitting of the training data

Lasso vs. Ridge Regularization

- Penalty is added to the loss function

(L1) Lasso: $\alpha \sum_i |w_i| \rightarrow$ sparse weights, mostly zeros

$$\nabla_{\vec{w}} \ell = -\frac{2}{m} X^T \vec{e} + 2\alpha * \text{sign}(\vec{w})$$

(L2) Ridge: $\alpha \sum_i w_i^2 \rightarrow$ small but distributed weights

$$\nabla_{\vec{w}} \ell = -\frac{2}{m} X^T \vec{e} + 2\alpha \vec{w} \rightarrow \text{gradient w/ ridge}$$