

Comparison with Gaussian Naïve Bayes

Gaussian Naïve Bayes vs. Logistic Regression

Generative $\frac{p(x|y)}{p(y)}$

Regression

Discriminative $\frac{p(y|x)}{1}$

Set of Gaussian
Naïve Bayes parameters
(feature variance
independent of class label)



Set of Logistic
Regression parameters

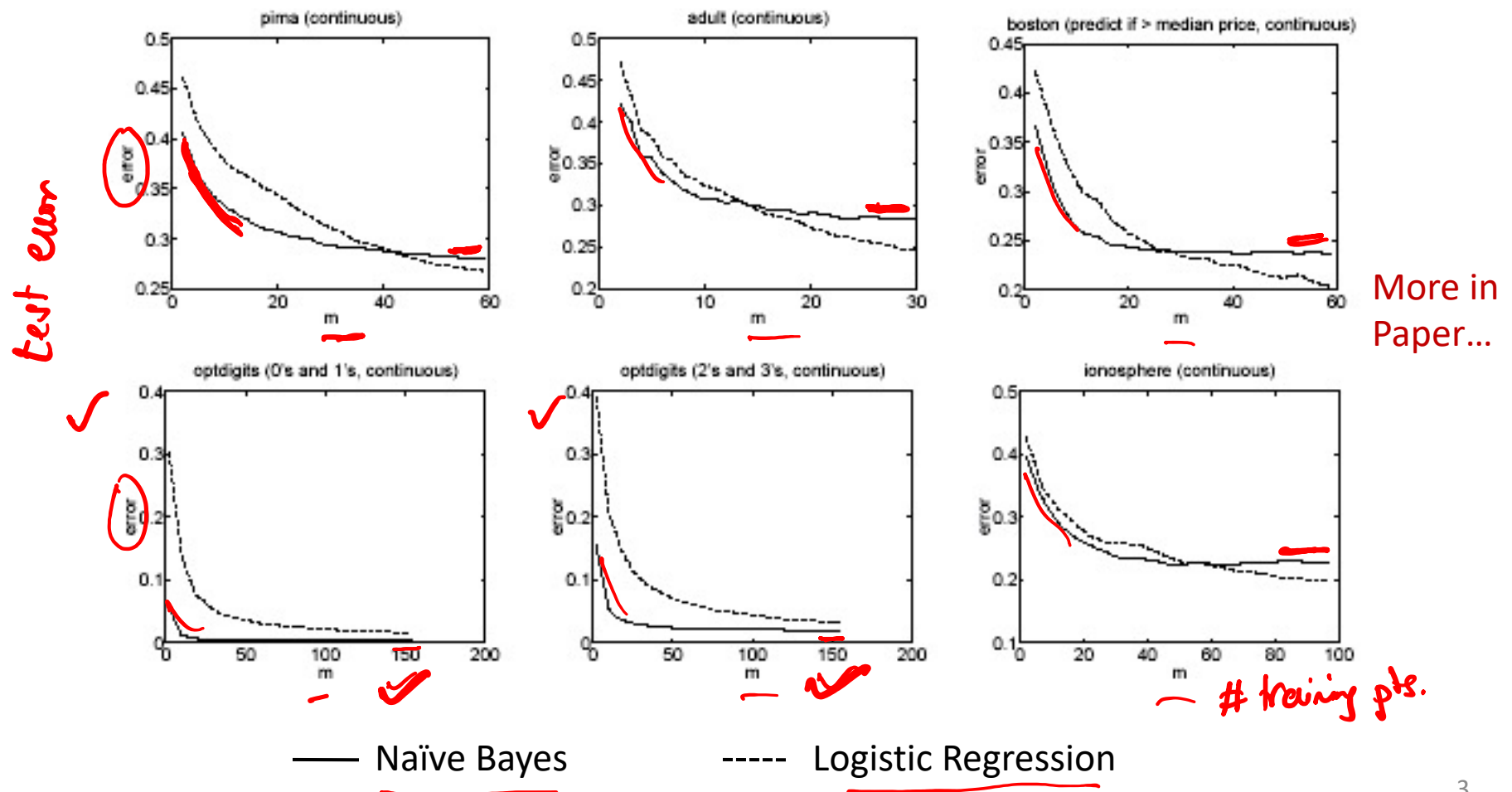
$\mu_y, \Sigma_y = \begin{bmatrix} \sigma_1^2 & 0 \\ 0 & \sigma_d^2 \end{bmatrix} \rightarrow \text{ind of } y \leftarrow$
then linear decision boundary

$w_0, w_1, \dots, w_d \leftarrow$

- Representation equivalence (both yield linear decision boundaries)
 - But only in a special case!!! (GNB with class-independent variances)
 - LR makes no assumptions about $P(X|Y)$ in learning!!!
 - Optimize different functions (MLE/MCLE) or (MAP/MCAP)! Obtain different solutions

Experimental Comparison (Ng-Jordan'01)

UCI Machine Learning Repository 15 datasets, 8 continuous features, 7 discrete features



Gaussian Naïve Bayes vs. Logistic Regression

- If conditional independence assumption holds, then GNB has lower large sample error
 $m \rightarrow \infty$
- If conditional independence assumption DOES NOT hold, then GNB has higher large sample error
 $m \rightarrow \infty$
 - But it converges faster (to its higher large sample error) as the parameter estimates are not coupled

What you should know

- LR is a linear classifier ✓
- LR optimized by maximizing conditional likelihood or conditional posterior
 - no closed-form solution
 - concave ! global optimum with gradient ascent
- Gaussian Naïve Bayes with class-independent variances representationally equivalent to LR
 - Solution differs because of objective (loss) function
- In general, NB and LR make different assumptions
 - NB: Features independent given class ! assumption on $P(\mathbf{X}|Y)$
 - LR: Functional form of $P(Y|\mathbf{X})$, no assumption on $P(\mathbf{X}|Y)$
- Convergence rates
 - GNB (usually) needs less data
 - LR (usually) gets to better solutions in the limit

Linear Regression

Aarti Singh

Machine Learning 10-315

Sept 20, 2021



MACHINE LEARNING DEPARTMENT



Supervised Learning Tasks

Classification

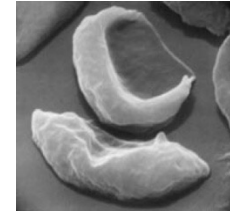


X = Document



Sports
Science
News

Y = Topic



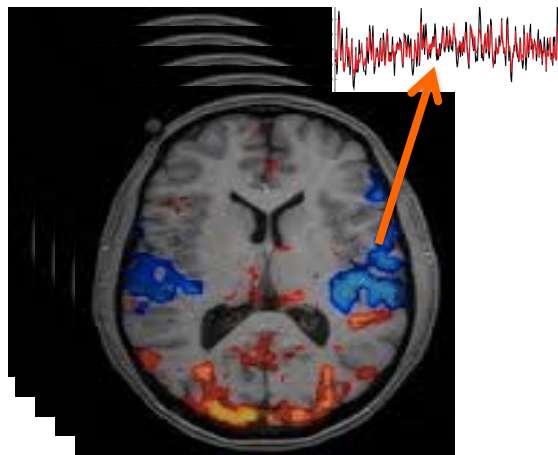
X = Cell Image



Anemic cell
Healthy cell

Y = Diagnosis

Regression



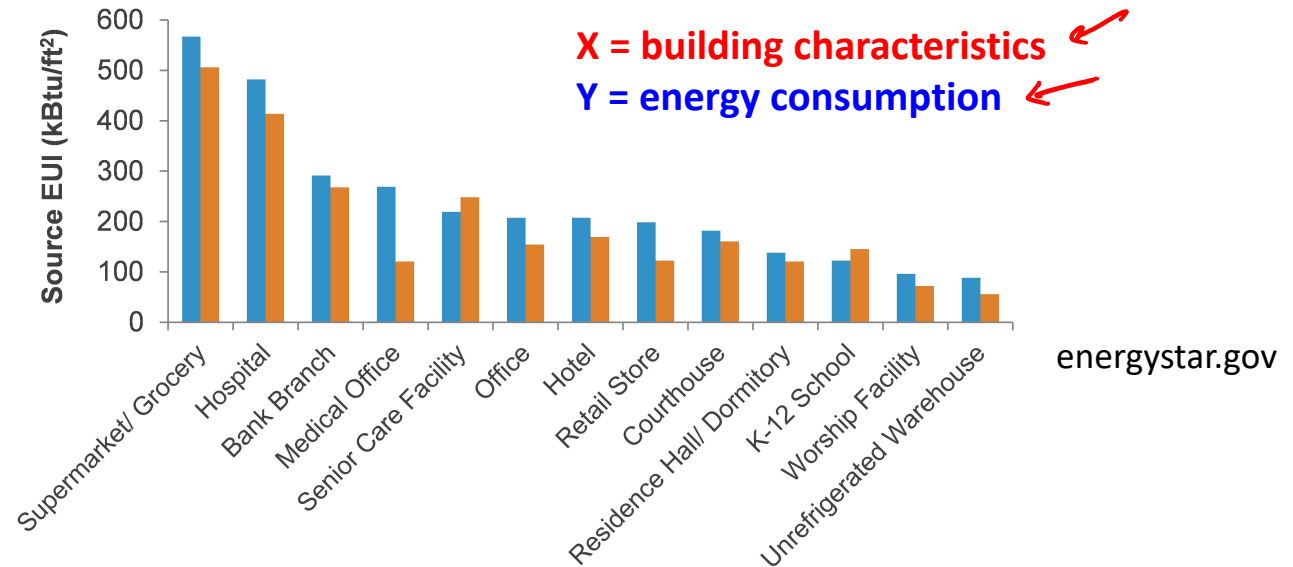
X = Brain Scan



Y = Age of a subject

Regression Tasks

Estimating Energy Usage



Estimating Contamination



Performance Measures

0/1

→ $\text{loss}(Y, f(X))$ - Measure of closeness between true label Y and prediction $f(X)$

Don't just want label of one test data (e.g. cell image), but any cell image $X \in \mathcal{X}$

$$(X, Y) \sim P_{XY}$$

Given a cell image drawn randomly from the collection of all cell images, how well does the predictor perform on average?

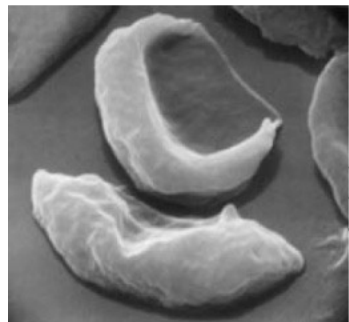
$$\text{Risk } R(f) \equiv \mathbb{E}_{XY} [\text{loss}(Y, f(X))]$$

0/1

$$E[1_{f(X) \neq Y}]$$
$$= P(f(X) \neq Y)$$

Performance Measures

Performance Measure: Risk $R(f) \equiv \mathbb{E}_{XY} [\text{loss}(Y, f(X))]$



➡ "Anemic cell"

$\text{loss}(Y, f(X))$

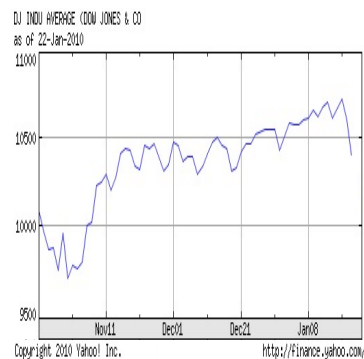
$$1_{\{f(X) \neq Y\}}$$

0/1 loss

Risk $R(f)$

$$P(f(X) \neq Y)$$

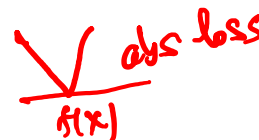
Probability of Error



➡ Share Price
"\$ 24.50"

$$(f(X) - Y)^2$$

square loss



$$\mathbb{E}[(f(X) - Y)^2]$$

Mean Square Error

MSE

Empirical Risk Minimization

Optimal predictor:

$$f^* = \arg \min_f \mathbb{E}_{\mathbf{x}, \mathbf{y}} [(f(X) - Y)^2]$$

Empirical Minimizer:

$$\hat{f}_n = \arg \min_{f \in \mathcal{F}} \underbrace{\frac{1}{n} \sum_{i=1}^n (f(X_i) - Y_i)^2}_{\text{Empirical mean}}$$

Law of Large Numbers:

$$\underbrace{\frac{1}{n} \sum_{i=1}^n [\text{loss}(Y_i, f(X_i))]}_{\text{Empirical mean}} \xrightarrow{n \rightarrow \infty} \mathbb{E}_{\mathbf{X}, \mathbf{Y}} [\text{loss}(Y, f(X))]$$

Restrict class of predictors



Optimal predictor:

$$f^* = \arg \min_f \mathbb{E}[(f(X) - Y)^2]$$

Empirical Minimizer:

$$\hat{f}_n = \arg \min_{f \in \mathcal{F}} \frac{1}{n} \sum_{i=1}^n (f(X_i) - Y_i)^2$$

Class of predictors

➤ Why?

- $\mathcal{F}$ - Class of Linear functions ✓ ←
- Class of Polynomial functions ✓
- Class of nonlinear functions ✓

Linear Regression

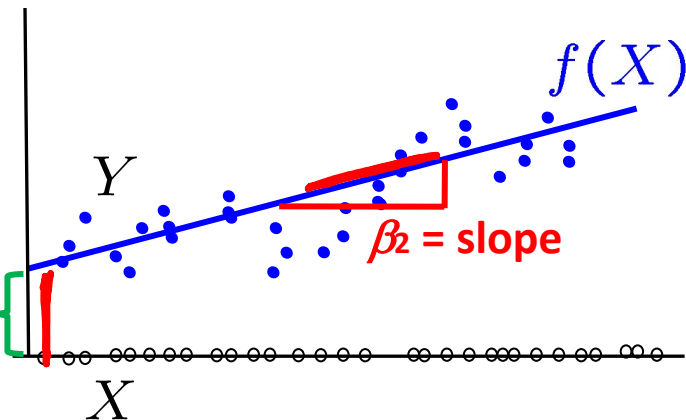
$$\hat{f}_n^L = \arg \min_{f \in \mathcal{F}_L} \frac{1}{n} \sum_{i=1}^n (f(X_i) - Y_i)^2 \quad \text{Least Squares Estimator}$$

$\mathcal{F}_L$ - Class of Linear functions

Uni-variate case:

$$f(X) = \beta_1 + \beta_2 X$$

β_1 - intercept



Multi-variate case:

p features

$$f(\underline{X}) = f(\underline{X}^{(1)}, \dots, \underline{X}^{(p)}) = \beta_1 \underline{X}^{(1)} + \beta_2 \underline{X}^{(2)} + \dots + \beta_p \underline{X}^{(p)}$$

$$= \underline{X} \beta \quad \text{where} \quad \underline{X} = [\underline{X}^{(1)} \dots \underline{X}^{(p)}], \quad \beta = [\beta_1 \dots \beta_p]^T$$

Linear Regression (Matrix-vector form)

$x_i^{(j)}$ - j^{th} feature of i^{th} data pt.

$x_i \leftarrow$ data point

$$\hat{f}_n^L = \arg \min_{f \in \mathcal{F}_L} \frac{1}{n} \sum_{i=1}^n (f(X_i) - Y_i)^2$$

$$\underline{f(X_i)} = \underline{X_i \beta}$$



$$\underline{\hat{\beta}} = \arg \min_{\underline{\beta}} \frac{1}{n} \sum_{i=1}^n (X_i \beta - Y_i)^2$$

$$\underline{\hat{f}_n^L(X)} = \underline{X \hat{\beta}}$$

$$= \arg \min_{\beta} \frac{1}{n} (\underline{A\beta - Y})^T (\underline{A\beta - Y}) = \frac{1}{n} \underline{z^T z}$$

$A\beta = \hat{Y}$
 $n \times p \quad p \times 1 \quad n \times 1$

$$\begin{matrix} n \times p \\ \rightarrow \\ A = \begin{bmatrix} X_1 \\ \vdots \\ X_n \end{bmatrix} = \begin{bmatrix} \rightarrow X_1^{(1)} & \dots & X_1^{(p)} \\ \vdots & \ddots & \vdots \\ \rightarrow X_n^{(1)} & \dots & X_n^{(p)} \end{bmatrix} \end{matrix}$$

Subjects

$$\begin{matrix} n \times 1 \\ Y = \begin{bmatrix} Y_1 \\ \vdots \\ Y_n \end{bmatrix} \end{matrix}$$

$z = A\beta - Y = \hat{Y} - Y$ $z^T z = \sum_{i=1}^n z_i^2$

Linear Regression

$$(A\beta)^T = \beta^T A^T$$

$$x^2 \cup$$

$$-x^2 \cap$$

$$x^T Q x$$

$$Q \text{ p.s.d.}$$

$$\hat{\beta} = \arg \min_{\beta} \frac{1}{n} (A\beta - Y)^T (A\beta - Y) = \arg \min_{\beta} J(\beta)$$

$$J(\beta) = (A\beta - Y)^T (A\beta - Y) \leftarrow (\beta^T A^T - Y^T) (A\beta - Y)$$

$$\rightarrow \beta^T A^T A \beta - Y^T A \beta - \beta^T A^T Y + Y^T Y$$

➤ Poll

Is the objective convex in β ?

Is $A^T A$ p.s.d.?

$$A = U \Lambda U^T$$

$n \times p$

$$A^T A = U \Lambda U^T U \Lambda U^T$$

$$= U \Lambda \Lambda^T U^T$$

$$= U \Lambda^2 U^T$$

A) Convex, quadratic in β

B) Non-convex, A may not be positive semi definite

C) Depends on conditioning (ratio of max:min eigenvalues) of $A^T A$

D) Convex, $A^T A$ is positive semi definite

$$\Lambda^2 = \begin{bmatrix} \lambda_1^2 & 0 \\ 0 & \lambda_r^2 \end{bmatrix}$$

Linear Regression

$$\beta = \begin{bmatrix} \beta_1 \\ \vdots \\ \beta_p \end{bmatrix}$$

$$\hat{\beta} = \arg \min_{\beta} \frac{1}{n} (\mathbf{A}\beta - \mathbf{Y})^T (\mathbf{A}\beta - \mathbf{Y}) = \arg \min_{\beta} J(\beta)$$

$$\mathbf{Y}^T \mathbf{A} \equiv \mathbf{Z}^T \quad \text{Rule 1}$$

$$J(\beta) = (\mathbf{A}\beta - \mathbf{Y})^T (\mathbf{A}\beta - \mathbf{Y}) \quad \leftarrow$$

$$= \beta^T \mathbf{A}^T \mathbf{A} \beta - \mathbf{Y}^T \mathbf{A} \beta - \beta^T \mathbf{A}^T \mathbf{Y} + \mathbf{Y}^T \mathbf{Y}$$

$$\frac{\partial J}{\partial \beta} = \mathbf{A}^T \mathbf{A} + \mathbf{A}^T \mathbf{A} \beta - 2 \mathbf{A}^T \mathbf{Y}$$

$$= 2 \mathbf{A}^T \mathbf{A} \beta - 2 \mathbf{A}^T \mathbf{Y}$$

$$\mathbf{A}^T \mathbf{A} \hat{\beta} = \mathbf{A}^T \mathbf{Y}$$

$$\left. \frac{\partial J(\beta)}{\partial \beta} \right|_{\hat{\beta}} = 0$$

$$\beta^T \mathbf{Z}$$

$$\frac{\partial \mathbf{Z}^T \beta}{\partial \beta} = \mathbf{Z}^T : \text{Rule 1}$$

$$\left. \frac{\partial \beta^T \mathbf{Q} \beta}{\partial \beta} = (\mathbf{Q} + \mathbf{Q}^T) \beta : \text{Rule 2} \right] \quad \begin{matrix} \mathbf{Q} = \mathbf{I} \\ \beta \text{ is } 1 \times 1 \end{matrix}$$

Linear regression solution satisfies

Normal Equations

$$A_{n \times p} = \begin{bmatrix} x_1^{(n)} & \dots & x_p^{(n)} \\ \vdots & & \vdots \\ x_1^{(p)} & \dots & x_p^{(p)} \end{bmatrix} \leftarrow$$

$$(\underbrace{A^T A}_{p \times p}) \underbrace{\hat{\beta}}_{p \times 1} = \underbrace{A^T Y}_{p \times 1}$$

$$A^T_{p \times n} \quad Y_{n \times 1} \quad Y = \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix}$$

If $(A^T A)$ is invertible,

$$\hat{\beta} = (A^T A)^{-1} A^T Y$$

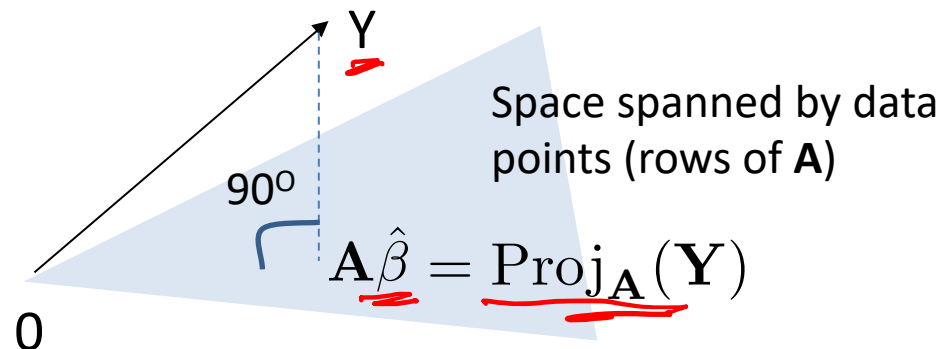
$$\hat{f}_n^L(X) = X \hat{\beta}$$

if $\underline{Y} = \underline{A} \underline{\beta}^*$

$\hat{f}_n^L(A) = A \beta^*$

$\hat{f}_n^L(A) = \underline{A} \hat{\beta} = \underbrace{A(A^T A)^{-1} A^T}_{\text{projection matrix}} Y$

Predicted labels for training points $\underline{A} \hat{\beta} = \text{Proj}_{\underline{A}}(\underline{Y})$



Linear regression solution satisfies Normal Equations

$$\underbrace{(\mathbf{A}^T \mathbf{A})}_{p \times p} \underbrace{\hat{\beta}}_{p \times 1} = \underbrace{\mathbf{A}^T \mathbf{Y}}_{p \times 1} \quad \leftarrow$$

If $(\mathbf{A}^T \mathbf{A})$ is invertible,

$$\hat{\beta} = (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T \mathbf{Y} \quad \leftarrow$$
$$\hat{f}_n^L(X) = X \hat{\beta}$$

Later: When is $(\mathbf{A}^T \mathbf{A})$ invertible ?

Now: What if $(\mathbf{A}^T \mathbf{A})$ is invertible but expensive (p very large)?

Gradient Descent

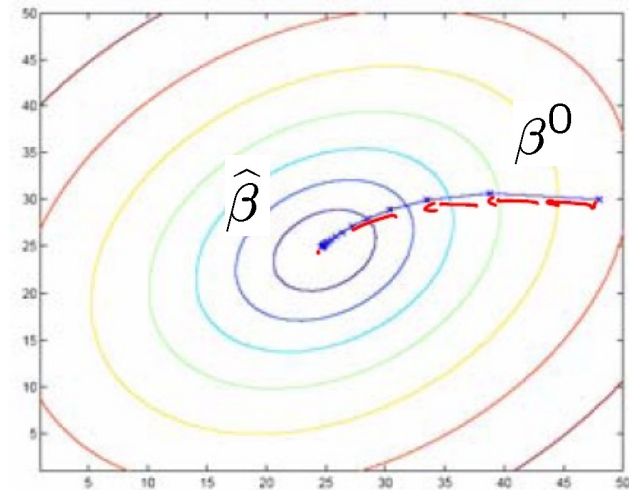
Even when $(\mathbf{A}^T \mathbf{A})$ is invertible, might be computationally expensive if $\mathbf{A}$ is huge.

$$\hat{\beta} = \arg \min_{\beta} \frac{1}{n} (\mathbf{A}\beta - \mathbf{Y})^T (\mathbf{A}\beta - \mathbf{Y}) = \arg \min_{\beta} \underline{\underline{J(\beta)}}$$

Since $J(\beta)$ is convex, move along negative of gradient

Initialize: β^0

$$\begin{aligned} \text{Update: } \beta^{t+1} &= \beta^t - \overset{\text{step size}}{\frac{\alpha}{2}} \frac{\partial J(\beta)}{\partial \beta} \bigg|_t \\ &= \beta^t - \alpha \underbrace{\mathbf{A}^T (\mathbf{A}\beta^t - \mathbf{Y})}_{0 \text{ if } \hat{\beta} = \beta^t} \end{aligned}$$



Stop: when some criterion met e.g. fixed # iterations, or $\frac{\partial J(\beta)}{\partial \beta} \bigg|_{\beta^t} < \epsilon$.

Least Square solution satisfies Normal Equations

$$\left. \frac{\partial J(\beta)}{\partial \beta} \right|_{\hat{\beta}} = 0 \quad \text{gives} \quad (\underset{p \times p}{\mathbf{A}^T \mathbf{A}}) \underset{p \times 1}{\hat{\beta}} = \underset{p \times 1}{\mathbf{A}^T \mathbf{Y}}$$

If $(\mathbf{A}^T \mathbf{A})$ is invertible,

1) If dimension p not too large, analytical solution:

$$\hat{\beta} = (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T \mathbf{Y} \quad \checkmark \quad \hat{f}_n^L(X) = X \hat{\beta} \quad \checkmark$$

2) If dimension p is large, computing inverse is expensive $O(p^3)$

Gradient descent since objective is convex ($\mathbf{A}^T \mathbf{A} \succeq 0$)

$$\begin{aligned} \beta^{t+1} &= \beta^t - \frac{\alpha}{2} \left. \frac{\partial J(\beta)}{\partial \beta} \right|_t \\ &= \beta^t - \alpha \mathbf{A}^T (\mathbf{A} \beta^t - \mathbf{Y}) \end{aligned} \quad \checkmark$$