

Regularized Linear Regression

Aarti Singh

Machine Learning 10-315

Sept 22, 2021



MACHINE LEARNING DEPARTMENT



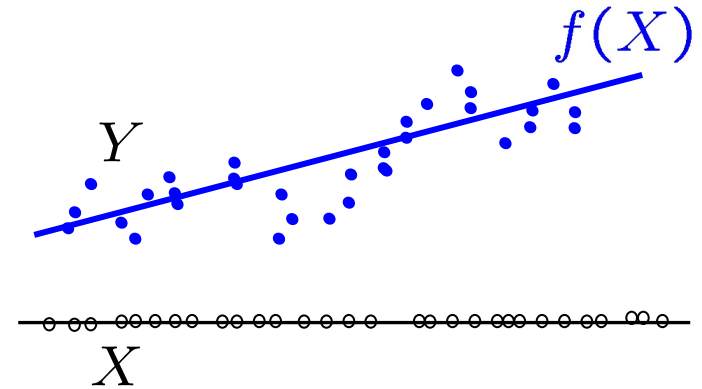
Linear Regression (Matrix-vector form)

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



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

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



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

$$\mathbf{A} = \begin{bmatrix} X_1 \\ \vdots \\ X_n \end{bmatrix} = \begin{bmatrix} X_1^{(1)} & \dots & X_1^{(p)} \\ \vdots & \ddots & \vdots \\ X_n^{(1)} & \dots & X_n^{(p)} \end{bmatrix} \quad \mathbf{Y} = \begin{bmatrix} Y_1 \\ \vdots \\ Y_n \end{bmatrix}$$

Least Square solution satisfies Normal Equations

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

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

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

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

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

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

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

Linear regression solution satisfies Normal Equations

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

$p \times p$ $p \times 1$ $p \times 1$

$\mathbf{A}^T \mathbf{A}$
 $p \times n$ $n \times p$
 $=$

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

Recall: Full rank matrices are invertible. What is rank of $(\mathbf{A}^T \mathbf{A})$?

If $\mathbf{A} = \mathbf{U} \mathbf{S} \mathbf{V}^T$, then
 $S - r \times r$

$$\begin{aligned}
 (\mathbf{V} \mathbf{S} \mathbf{U}^T) \hat{\beta} &= \mathbf{V} \mathbf{S} \mathbf{U}^T \mathbf{Y} \\
 \mathbf{V}^T \mathbf{V} \mathbf{S}^2 \mathbf{V}^T \hat{\beta} &= \mathbf{V}^T \mathbf{V} \mathbf{S} \mathbf{U}^T \mathbf{Y} \\
 \mathbf{S}^{-1} \mathbf{S}^2 \mathbf{V}^T \hat{\beta} &= \mathbf{S}^T \mathbf{U}^T \mathbf{Y} \\
 \mathbf{S} \mathbf{V}^T \hat{\beta} &= \mathbf{U}^T \mathbf{Y}
 \end{aligned}$$

normal equations $(\mathbf{S} \mathbf{V}^T) \hat{\beta} = (\mathbf{U}^T \mathbf{Y})$
 $r \times p$ $p \times 1$ $r \times 1$

r equations in p unknowns.

Under-determined if $r < p$, hence no unique solution.

Regularized Least Squares

What if $(\mathbf{A}^T \mathbf{A})$ is not invertible? ✓

$$r \leq \min(n, p)$$

r equations, p unknowns – underdetermined system of linear equations
many feasible solutions

Need to constrain solution further

e.g. bias solution to “small” values of β (small changes in input don’t translate to large changes in output)

$$\hat{\beta}_{\text{MAP}} = \arg \min_{\beta} \sum_{i=1}^n (Y_i - X_i \beta)^2 + \lambda \|\beta\|_2^2$$

Ridge Regression
(l2 penalty)

$$= \arg \min_{\beta} (\mathbf{A}\beta - \mathbf{Y})^T (\mathbf{A}\beta - \mathbf{Y}) + \lambda \|\beta\|_2^2$$

$$\lambda \geq 0$$

$$\hat{\beta}_{\text{MAP}} = (\mathbf{A}^T \mathbf{A} + \lambda \mathbf{I})^{-1} \mathbf{A}^T \mathbf{Y}$$

Regularized Least Squares

$$\beta^T A^T A \beta \leftarrow$$

$$\frac{\partial \beta^T M \beta}{\partial \beta} = (M + M^T) \beta$$

$$\hat{\beta}_{\text{MAP}} = \arg \min_{\beta} \sum_{i=1}^n (Y_i - X_i \beta)^2 + \lambda \|\beta\|_2^2$$

Ridge Regression
(l2 penalty)

$$= \arg \min_{\beta} \frac{1}{n} (\mathbf{A}\beta - \mathbf{Y})^T (\mathbf{A}\beta - \mathbf{Y}) + \lambda \|\beta\|_2^2, \quad \lambda \geq 0$$

$$\frac{\partial J(\beta)}{\partial \beta} = 2 \mathbf{A}^T \mathbf{A} \beta - 2 \mathbf{A}^T \mathbf{Y} + \lambda \frac{\partial \|\beta\|_2^2}{\partial \beta}$$

$$\|\beta\|_2^2 = \beta^T \beta \leftarrow$$

$$= \sum_i \beta^{(i)2}$$

$$= 2 (\mathbf{A}^T \mathbf{A} + \lambda \mathbf{I}) \beta - 2 \mathbf{A}^T \mathbf{Y} \quad = 2 \lambda \beta$$

$$= 0 \Rightarrow \hat{\beta} = (\mathbf{A}^T \mathbf{A} + \lambda \mathbf{I})^{-1} \mathbf{A}^T \mathbf{Y}$$

$$\text{eval}(\mathbf{A}^T \mathbf{A} + \lambda \mathbf{I}) = \text{eval}(\mathbf{A}^T \mathbf{A}) + \lambda$$

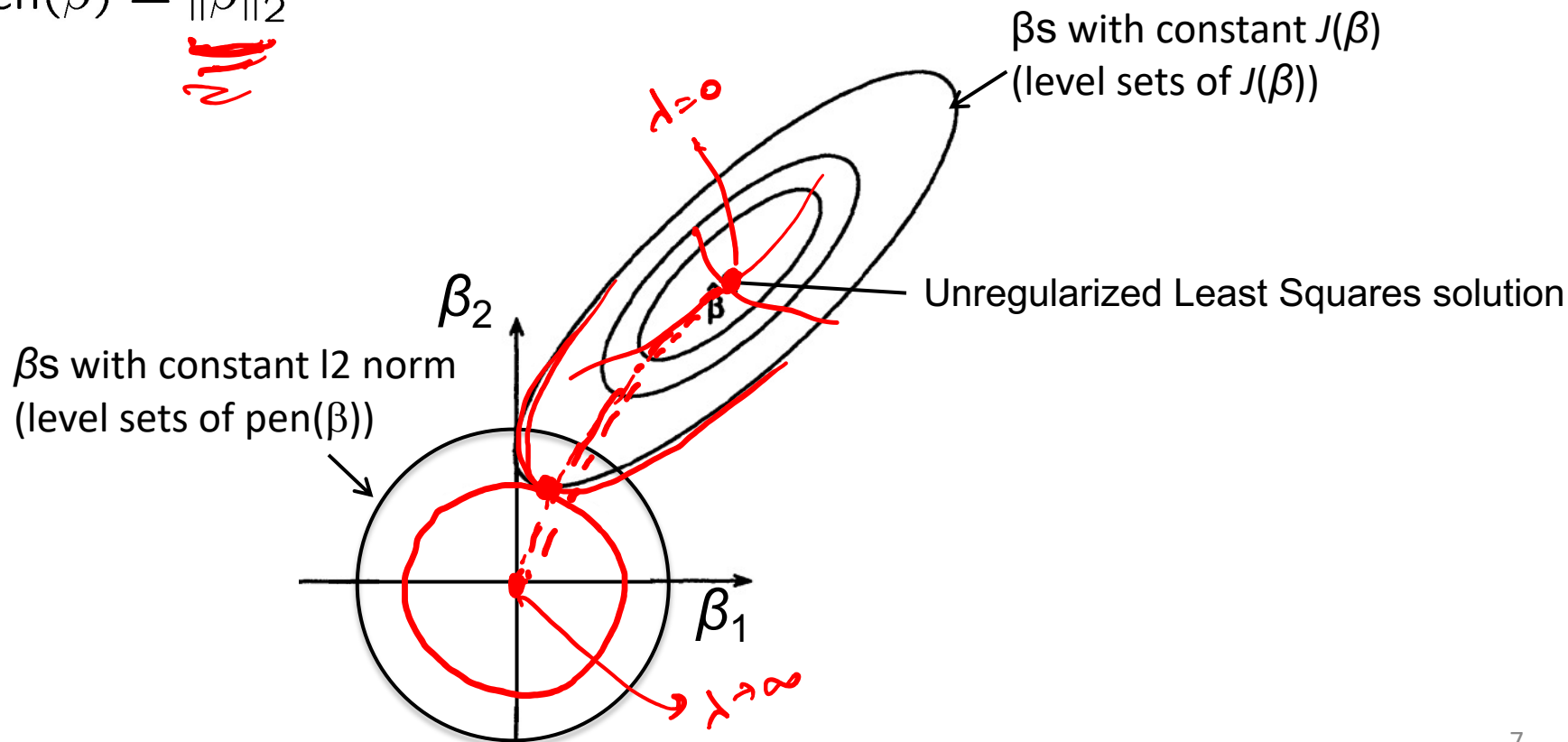
$p \times p$
Is $(\mathbf{A}^T \mathbf{A} + \lambda \mathbf{I})$ invertible?

Understanding regularized Least Squares

$$\min_{\beta} \underbrace{(\mathbf{A}\beta - \mathbf{Y})^T (\mathbf{A}\beta - \mathbf{Y})}_{J(\beta)} + \lambda \underbrace{\text{pen}(\beta)}_{\text{penalty}} = \min_{\beta} J(\beta) + \lambda \text{pen}(\beta)$$

Ridge Regression:

$$\text{pen}(\beta) = \|\beta\|_2^2$$



Regularized Least Squares

What if $(\mathbf{A}^T \mathbf{A})$ is not invertible ?

r equations , p unknowns – underdetermined system of linear equations
many feasible solutions

Need to constrain solution further

e.g. bias solution to “small” values of β (small changes in input don’t translate to large changes in output)

$$\hat{\beta}_{\text{MAP}} = \arg \min_{\beta} \sum_{i=1}^n (Y_i - X_i \beta)^2 + \lambda \|\beta\|_2^2$$

Ridge Regression
(l2 penalty)

$$\hat{\beta}_{\text{MAP}} = \arg \min_{\beta} \sum_{i=1}^n (Y_i - X_i \beta)^2 + \lambda \|\beta\|_1$$

Lasso
(l1 penalty)

$\lambda \geq 0$

Many β can be zero – many inputs are irrelevant to prediction in high-dimensional settings (typically intercept term not penalized)

Regularized Least Squares

What if $(\mathbf{A}^T \mathbf{A})$ is not invertible ?

r equations , p unknowns – underdetermined system of linear equations
many feasible solutions

Need to constrain solution further

e.g. bias solution to “small” values of β (small changes in input don’t translate to large changes in output)

$$\hat{\beta}_{\text{MAP}} = \arg \min_{\beta} \sum_{i=1}^n (Y_i - X_i \beta)^2 + \lambda \|\beta\|_2^2$$

Ridge Regression
(l2 penalty)


$$\hat{\beta}_{\text{MAP}} = \arg \min_{\beta} \sum_{i=1}^n (Y_i - X_i \beta)^2 + \lambda \|\beta\|_1$$

$$= \sum_j |\beta_j|$$

Lasso
(l1 penalty) $\lambda \geq 0$

No closed form solution, but can optimize using sub-gradient descent (packages available)

Ridge Regression vs Lasso

$$\min_{\beta} (\mathbf{A}\beta - \mathbf{Y})^T (\mathbf{A}\beta - \mathbf{Y}) + \lambda \text{pen}(\beta) = \min_{\beta} J(\beta) + \lambda \text{pen}(\beta)$$

Ridge Regression:

$$\text{pen}(\beta) = \|\beta\|_2^2$$

= Const

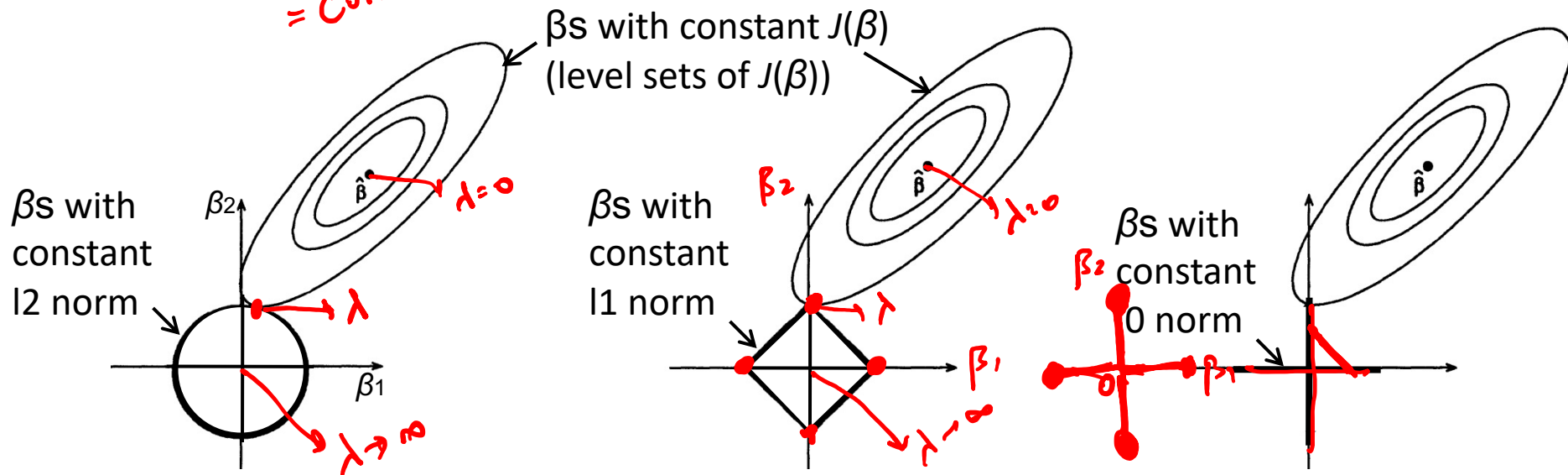
Lasso:

$$\text{pen}(\beta) = \|\beta\|_1$$

= Const.

Ideally l0 penalty,
but optimization
becomes non-convex

=



Lasso (l1 penalty) results in sparse solutions – vector with more zero coordinates
Good for high-dimensional problems – don't have to store all coordinates, interpretable solution!

$n < p$

Matlab example

```
clear all  
close all
```

```
n = 80;    % datapoints  
p = 100;  % features  
k = 10;    % non-zero features
```

```
rng(20);  
X = randn(n,p);  
weights = zeros(p,1);  
weights(1:k) = randn(k,1)+10;  
noise = randn(n,1) * 0.5;  
Y = X*weights + noise;
```

```
Xtest = randn(n,p);  
noise = randn(n,1) * 0.5;  
Ytest = Xtest*weights + noise;
```

```
lassoWeights = lasso(X,Y,'Lambda',1,  
'Alpha', 1.0);  
Ylasso = Xtest*lassoWeights;  
norm(Ytest-Ylasso)
```

```
ridgeWeights = lasso(X,Y,'Lambda',1,  
'Alpha', 0.0001);  
Yridge = Xtest*ridgeWeights;  
norm(Ytest-Yridge)
```

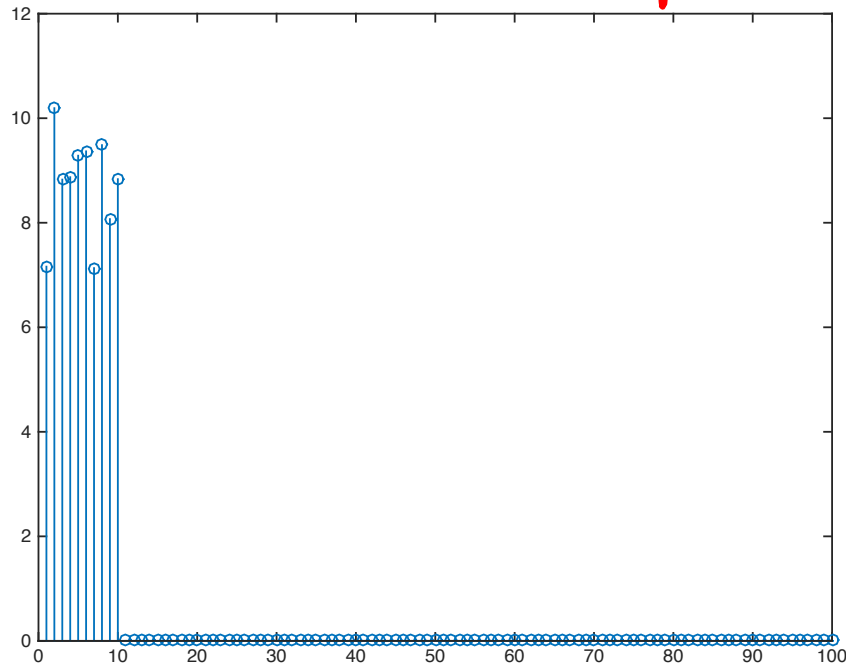
```
stem(lassoWeights)  
pause  
stem(ridgeWeights)
```

Matlab example

Test MSE = 33.7997

Lasso Coefficients

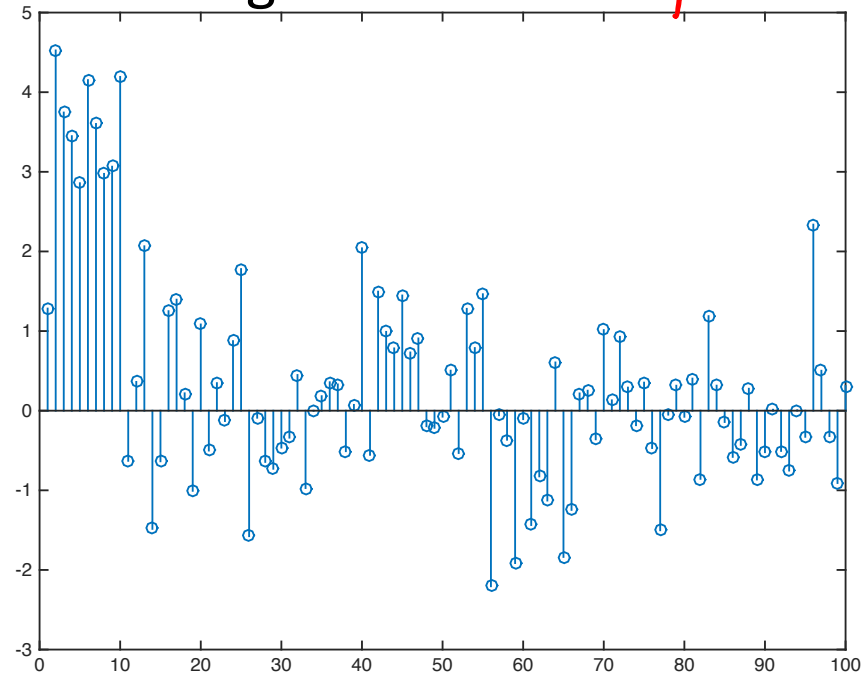
β^1



Test MSE = 185.9948

Ridge Coefficients

β^1



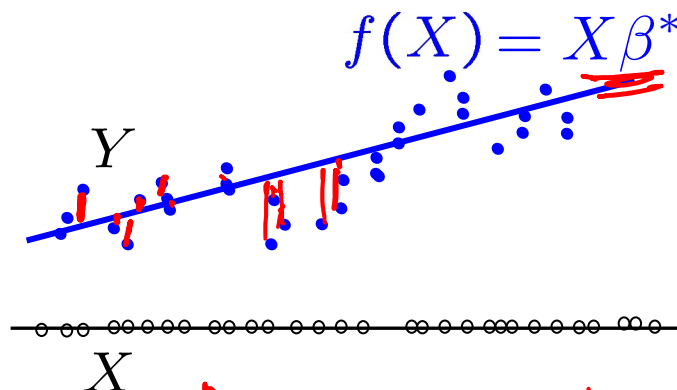
$\sum_i (y_i - x_i \beta)^2$ Least Squares and M(C)LE

Intuition: Signal plus (zero-mean) Noise model

$$Y = \underline{f^*(X)} + \underline{\epsilon} = \underline{X\beta^*} + \underline{\epsilon}$$

$$\epsilon \sim \underline{\mathcal{N}(0, \sigma^2 \mathbf{I})} \quad Y \sim \underline{\mathcal{N}(X\beta^*, \sigma^2 \mathbf{I})}$$

$$\hat{\beta}_{\text{MLE}} = \arg \max_{\beta} \underbrace{\log p(\{Y_i\}_{i=1}^n | \beta, \sigma^2, \{X_i\}_{i=1}^n)}_{\text{Conditional log likelihood}}$$



$$\prod_{i=1}^n p(Y_i | X_i; \beta, \sigma^2)$$

$$= \arg \max_{\beta} \sum_{i=1}^n \log p(Y_i | X_i; \beta, \sigma^2)$$

$$= \arg \max_{\beta} \sum_{i=1}^n \log \frac{e^{-\frac{(Y_i - X_i \beta^*)^2}{2\sigma^2}}}{\sqrt{2\pi\sigma^2}} = \arg \max_{\beta} - \sum_{i=1}^n \frac{(Y_i - X_i \beta^*)^2}{2\sigma^2}$$

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

Least Square Estimate is same as Maximum
Conditional Likelihood Estimate under a Gaussian noise model !