

Regularized Linear Regression

Mahdi Roozbahani
Georgia Tech

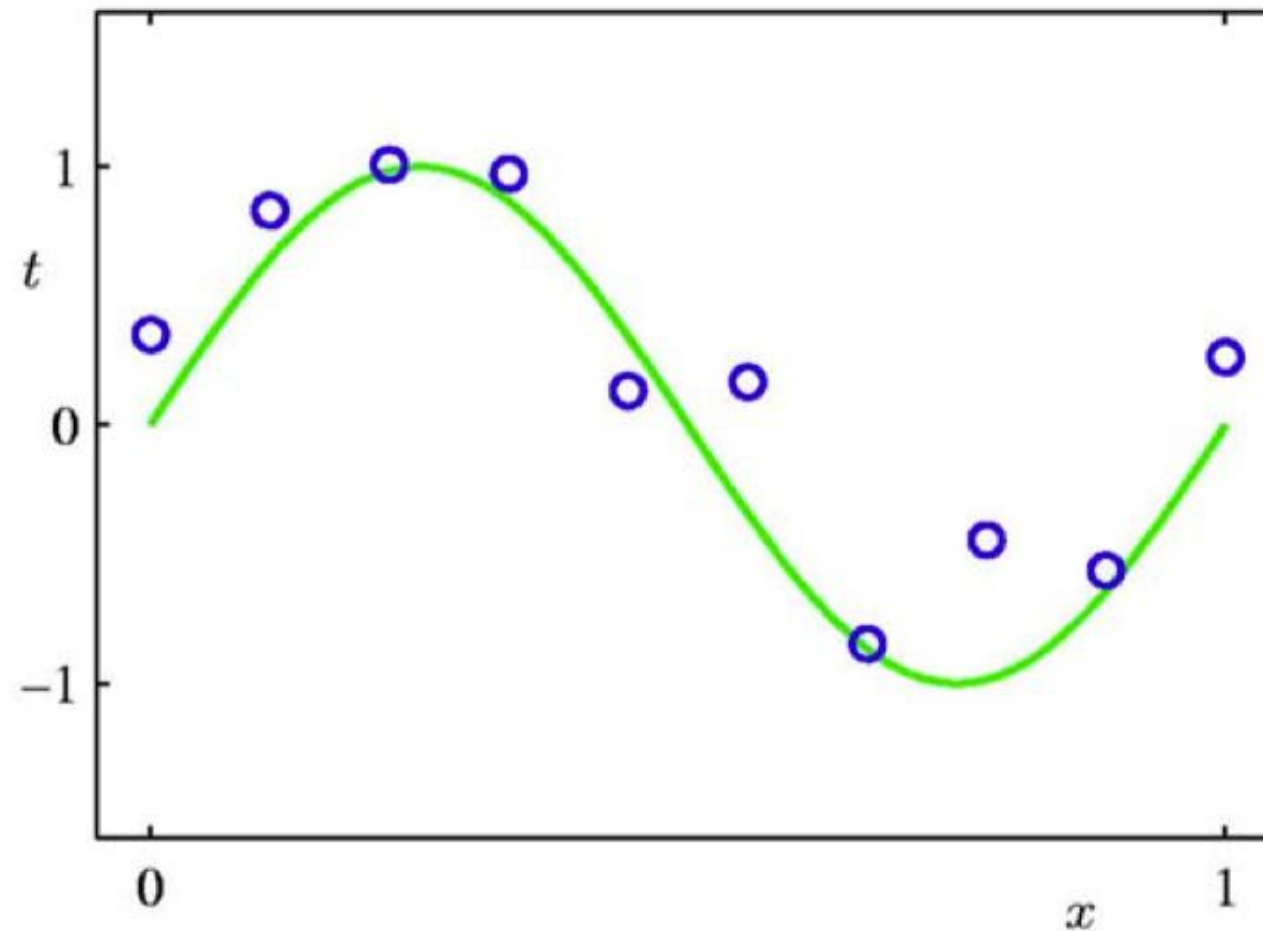
A photograph of a wooden floor with a mattress placed on it. The mattress is shaped into a complex, winding, and convoluted path, resembling a maze or a highly irregular line. This visual metaphor is used to represent overfitting, where a model is too closely tailored to a specific dataset and thus fails to generalize to new data. The text "THE BEST WAY TO EXPLAIN OVERFITTING" is overlaid in large, bold, white capital letters with a black outline at the bottom of the image.

**THE BEST WAY TO
EXPLAIN OVERFITTING**

Outline

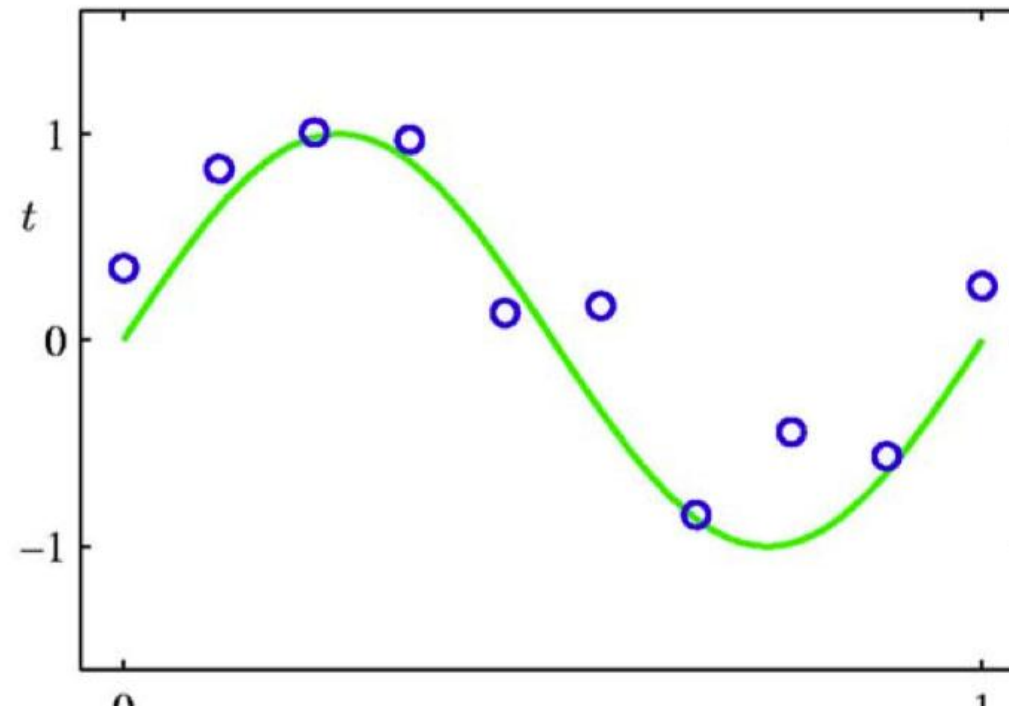
- Overfitting and regularized learning 
- Ridge regression
- Lasso regression
- Determining regularization strength

Regression: Recap



- Suppose we are given a training set of N observations
- Regression problem is to estimate $y(x)$ from this data

Regression: Recap



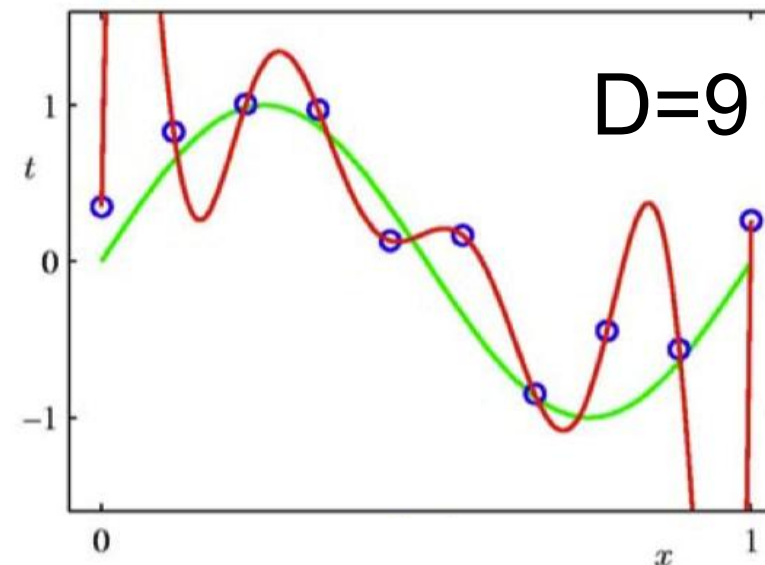
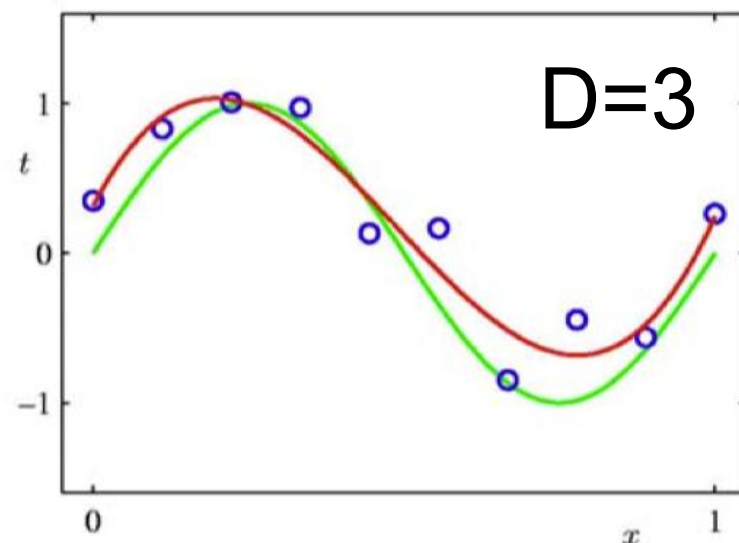
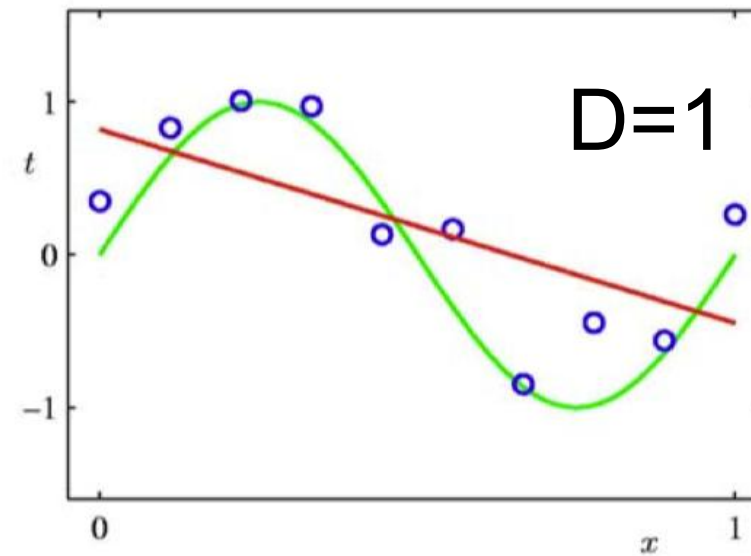
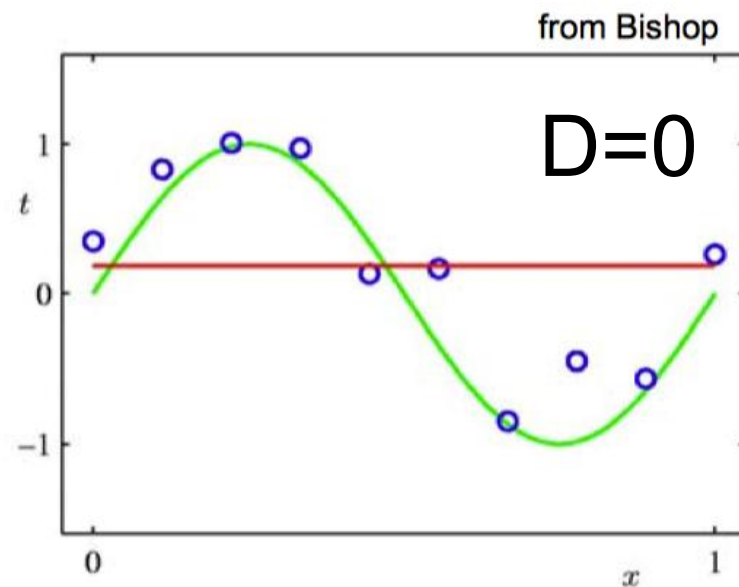
- Want to fit a polynomial regression model

$$y = \theta_0 + \theta_1 x + \theta_2 x^2 + \dots + \theta_d x^d + \epsilon$$

- $z = \{1, x, x^2, \dots, x^d\} \in R^d$ and $\theta = (\theta_0, \theta_1, \theta_2, \dots, \theta_d)^T$

$$y = z\theta$$

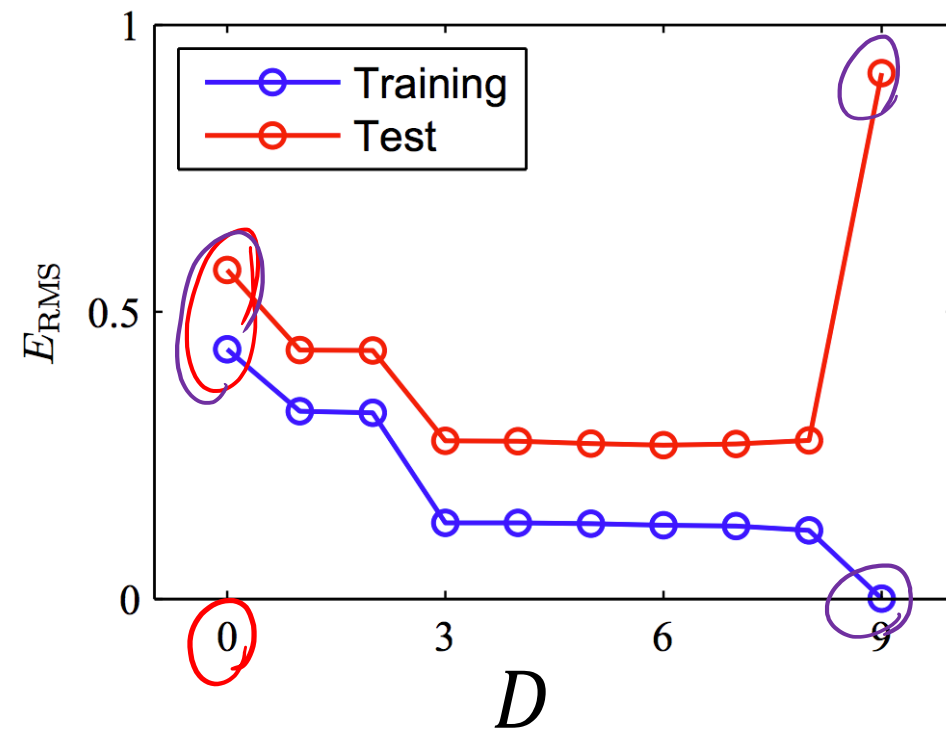
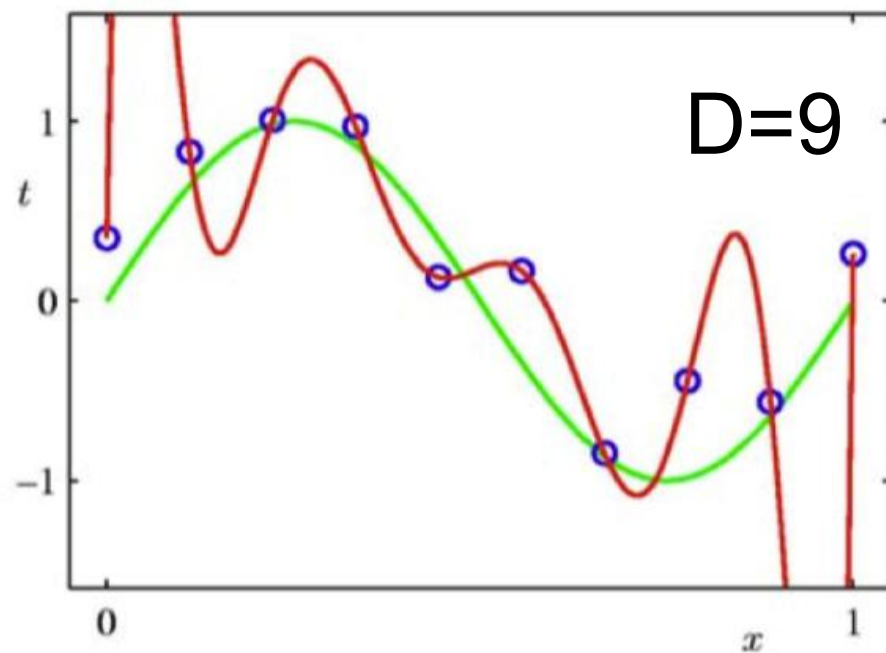
Which One is Better?



- Can we increase the maximal polynomial degree to very large, such that the curve passes through all training points?

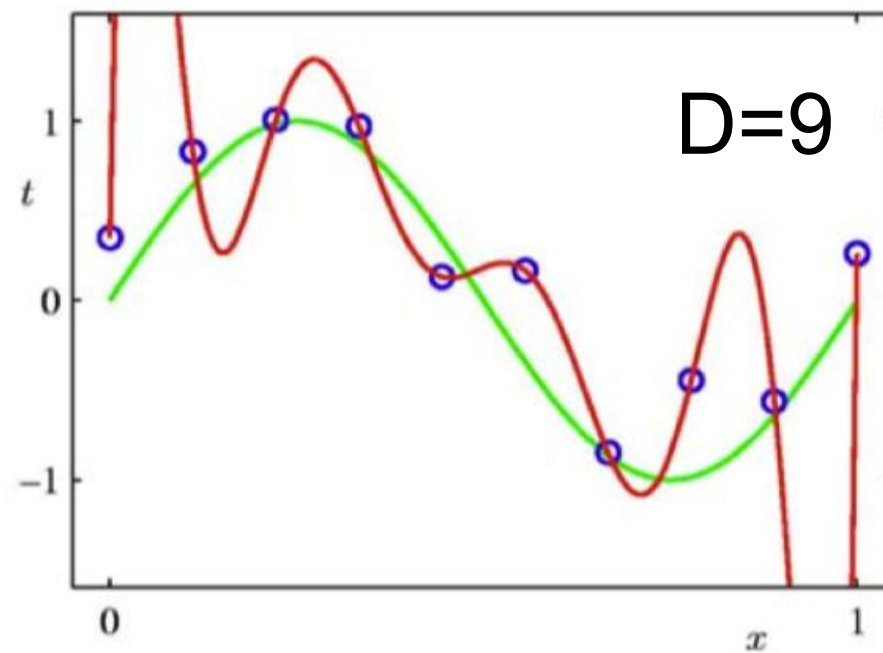
No, this can lead to **overfitting**!

The Overfitting Problem



- The training error is very low, but the error on test set is large.
- The model captures not only patterns but also noisy nuisances in the training data.

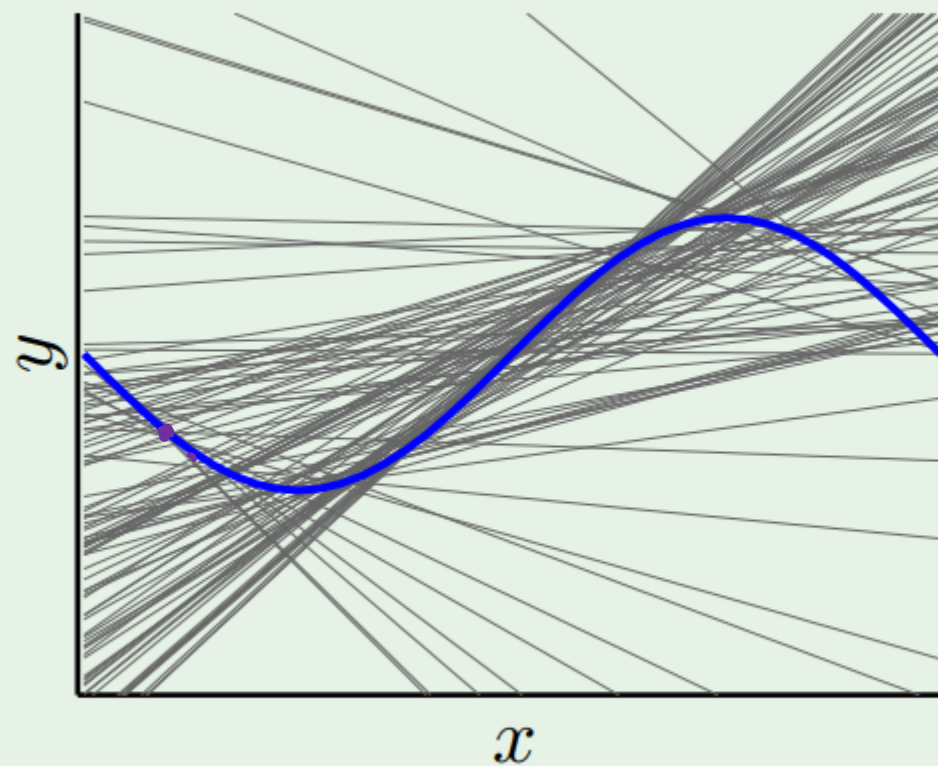
The Overfitting Problem



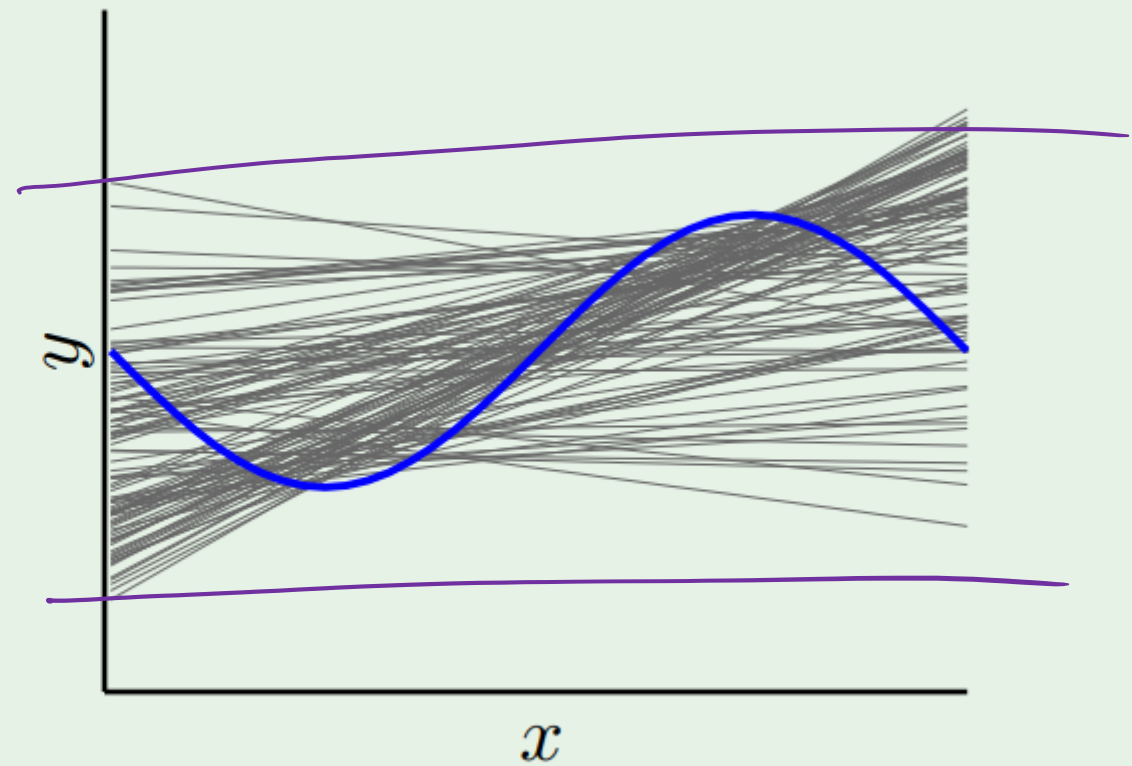
- In regression, overfitting is often associated with large Weights (**severe oscillation**)
- How can we address overfitting?

Regularization

(smart way to cure overfitting disease)



without regularization



with regularization

Put a brake on fitting

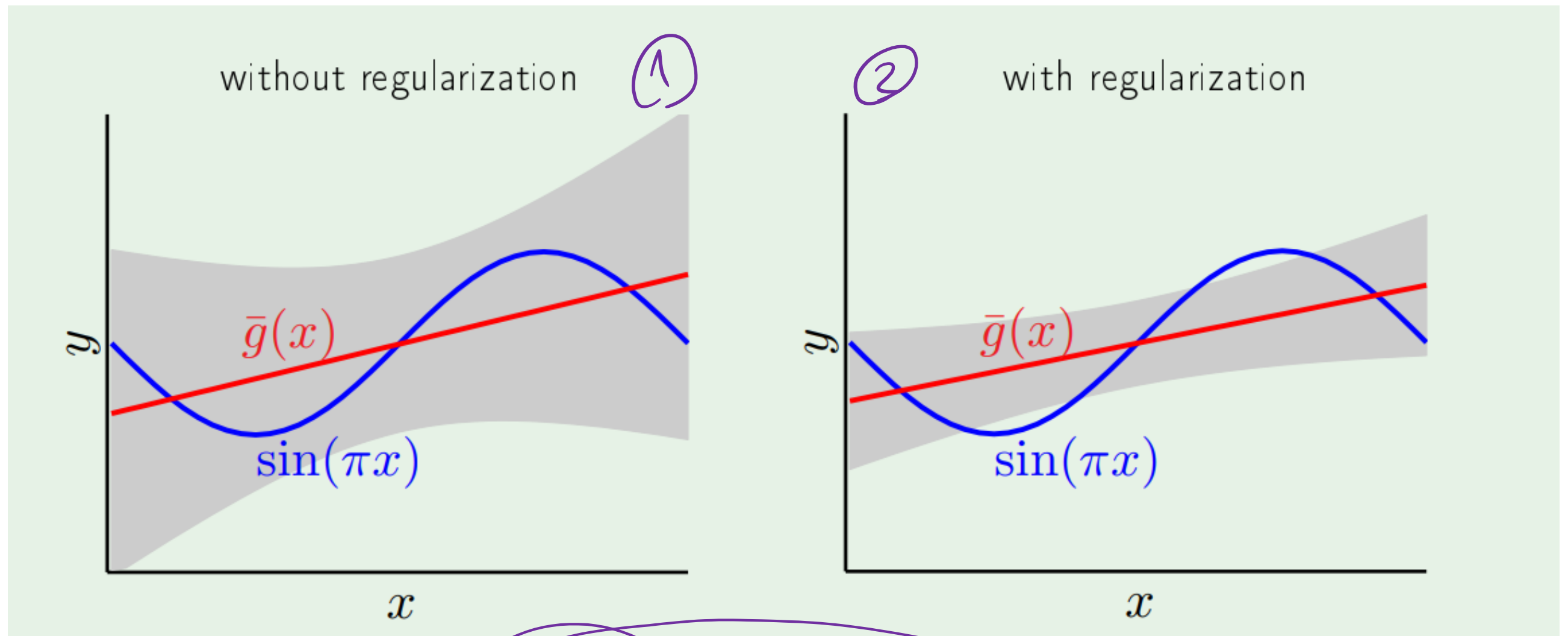


Fit a linear line on sinusoidal with just two data points

Who is the winner?

$$E(\theta) = \text{bias}^2 + \text{variance}$$

$\bar{g}(x)$: average over all lines



bias=0.21; var=1.69

bias=0.23; var=0.33

Polynomial Model

Want to fit a polynomial regression model

$$y = \theta_0 + \theta_1 x + \theta_2 x^2 + \cdots + \theta_d x^d + \epsilon$$

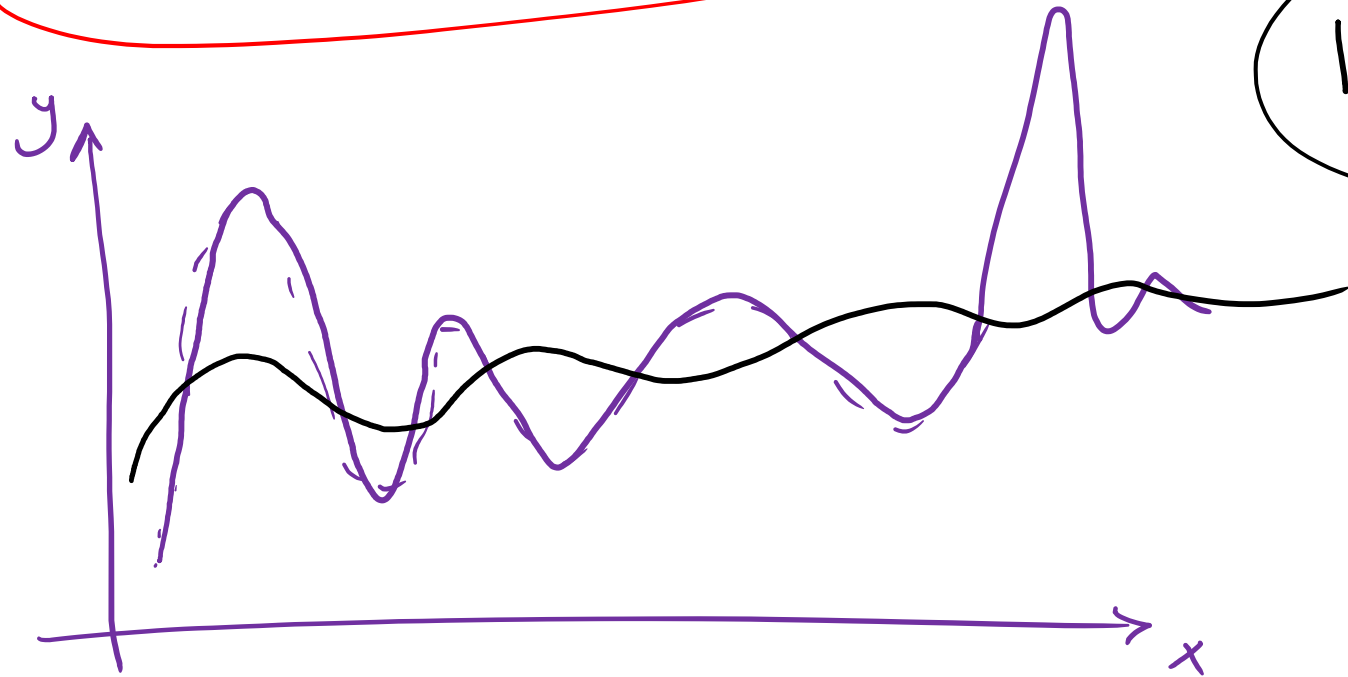
Let's rewrite it as:

$$y = \theta_0 + \theta_1 z_1 + \theta_2 z_2 + \cdots + \theta_d z_d + \epsilon = \mathbf{z}\boldsymbol{\theta}$$

$$y = \underbrace{\theta_0}_{=} + \underbrace{\theta_1}_{=} x + \underbrace{\theta_2}_{=} x^2 + \dots + \underbrace{\theta_d}_{=} x^d \rightarrow \text{LCF}$$

hyper parameter

$$|\theta_0| + |\theta_1| + \dots + |\theta_d| \leq 100$$



$$X = \begin{bmatrix} \theta_1 & \theta_2 & \dots & \theta_d \\ x_1 & x_2 & \dots & x_d \end{bmatrix} \quad y = \begin{bmatrix} y^{(1)} \\ y^{(2)} \\ \dots \\ y^{(n)} \end{bmatrix}$$

$$\theta_0 + \theta_1 x_1 + \theta_2 x_2 + \dots + \theta_d x_d \approx y^{(i)}$$

Regularizing is just constraining the weights (θ)

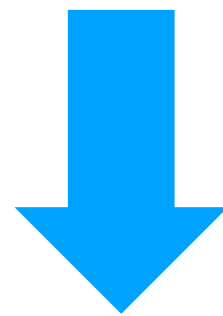
For example: let's do a **hard** constraining

$$y = \theta_0 + \theta_1 z_1 + \theta_2 z_2 + \cdots + \theta_d z_d$$

$$x = \begin{bmatrix} z_1 & z_2 & z_3 & \dots & z_d \\ h & h^2 & h^3 & \dots & h^d \\ x_1 & x_2 & & & x_d \end{bmatrix} \quad x_d = h^d$$

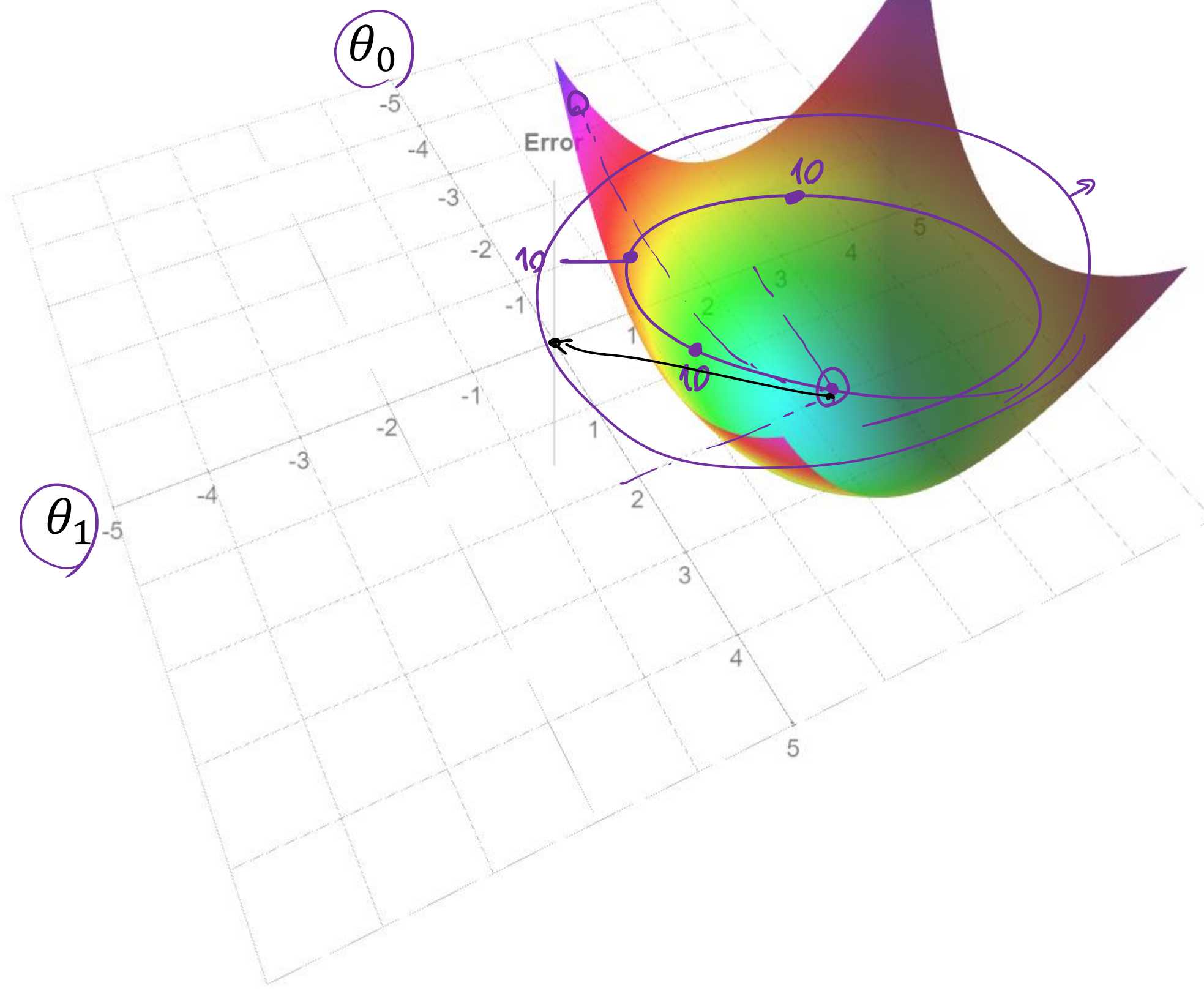
subject to

$$\theta_d = 0 \text{ for } d > 2$$

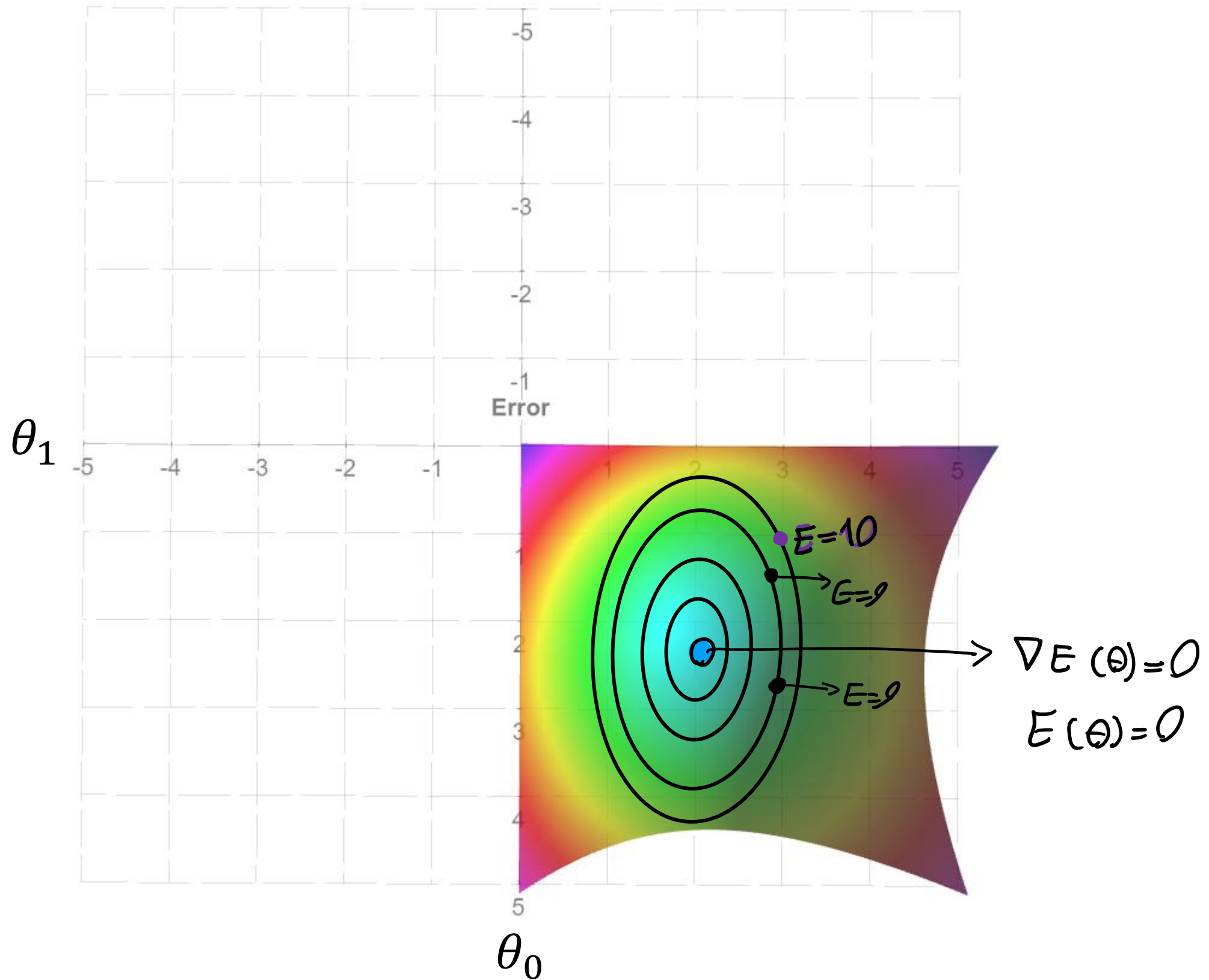


$$y = \theta_0 + \theta_1 z_1 + \theta_2 z_2 + 0 + \cdots + 0$$

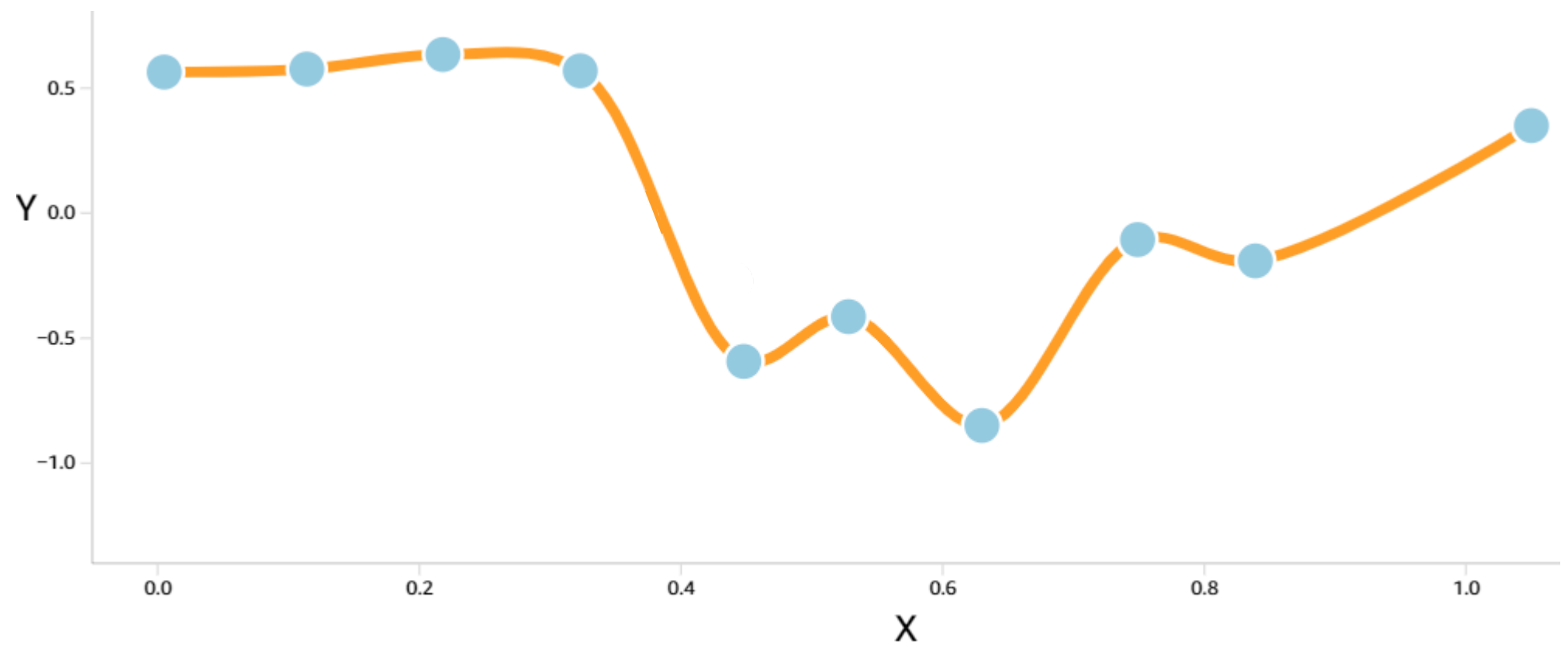
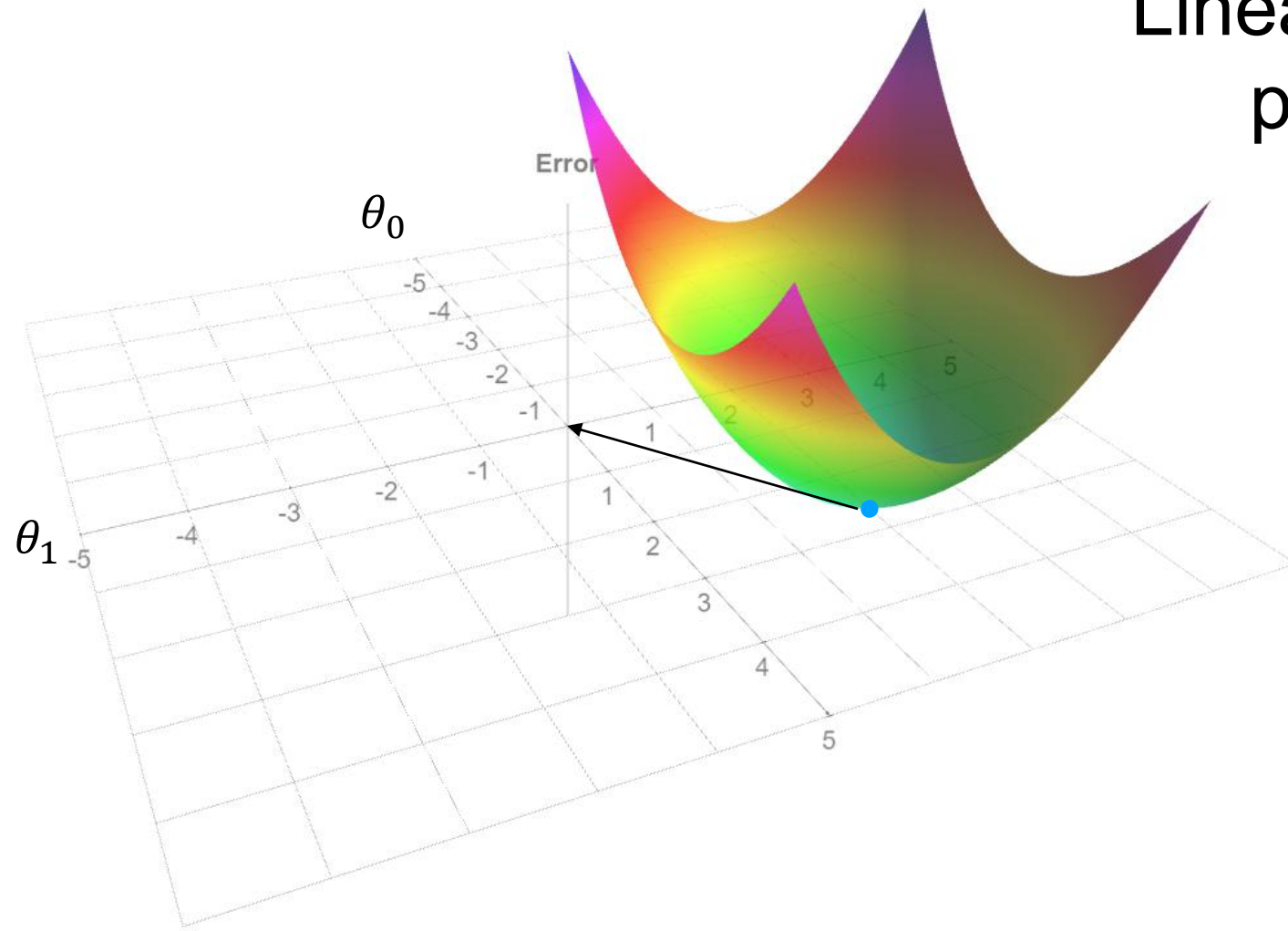
$$E(\theta) = \frac{1}{N} \sum_{i=1}^n (y^{(i)} - z^{(i)} \theta)^2 \xrightarrow{\text{min}} \text{Convex}$$

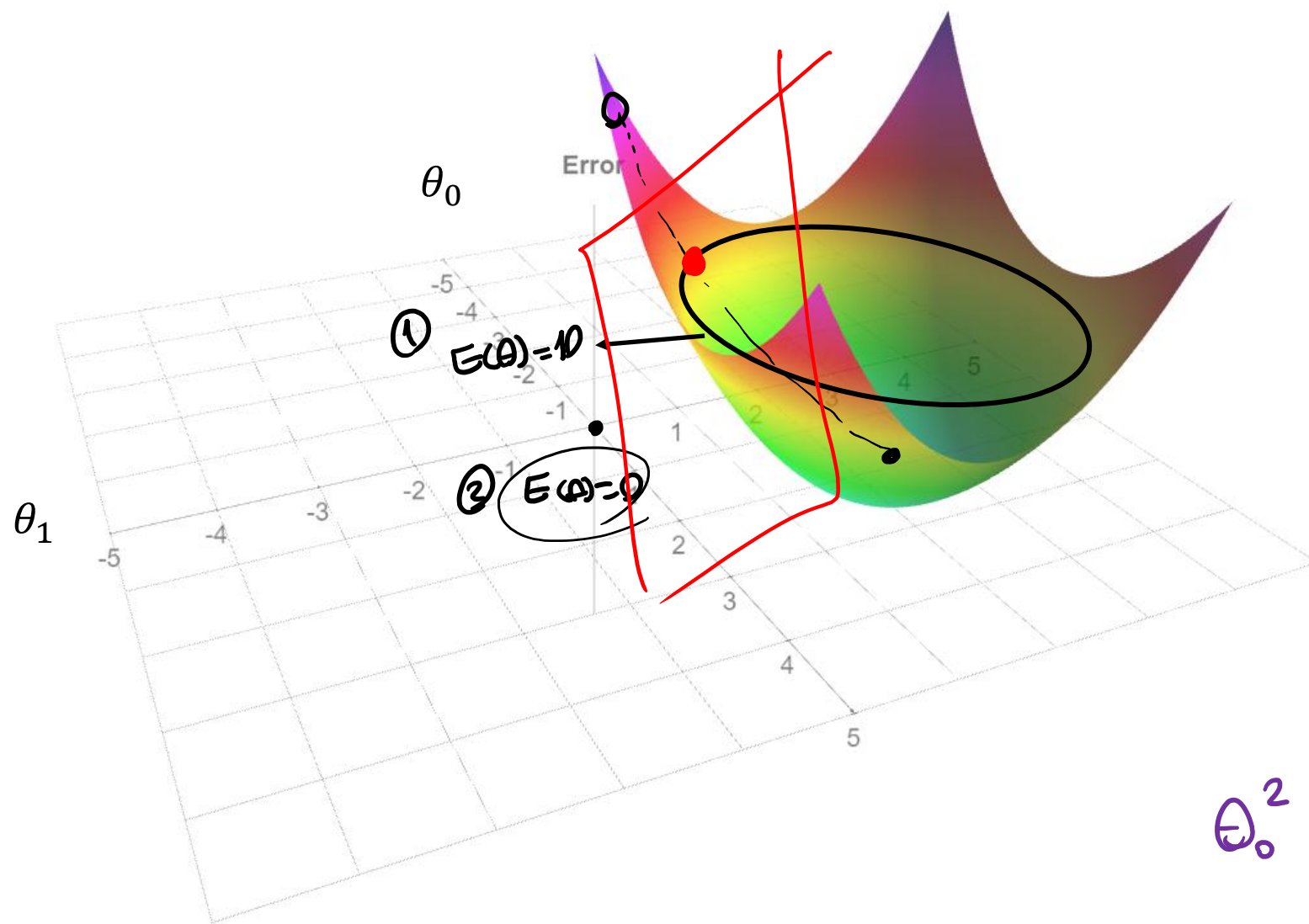


Project the same graph on x-y using contour plot



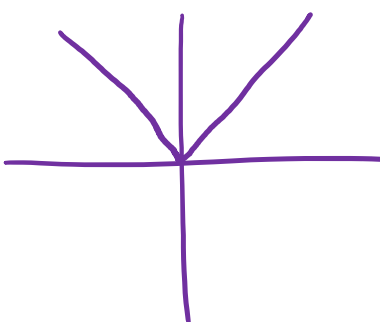
Linear regression with a very high polynomial degree solution



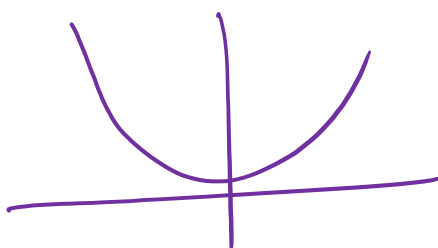


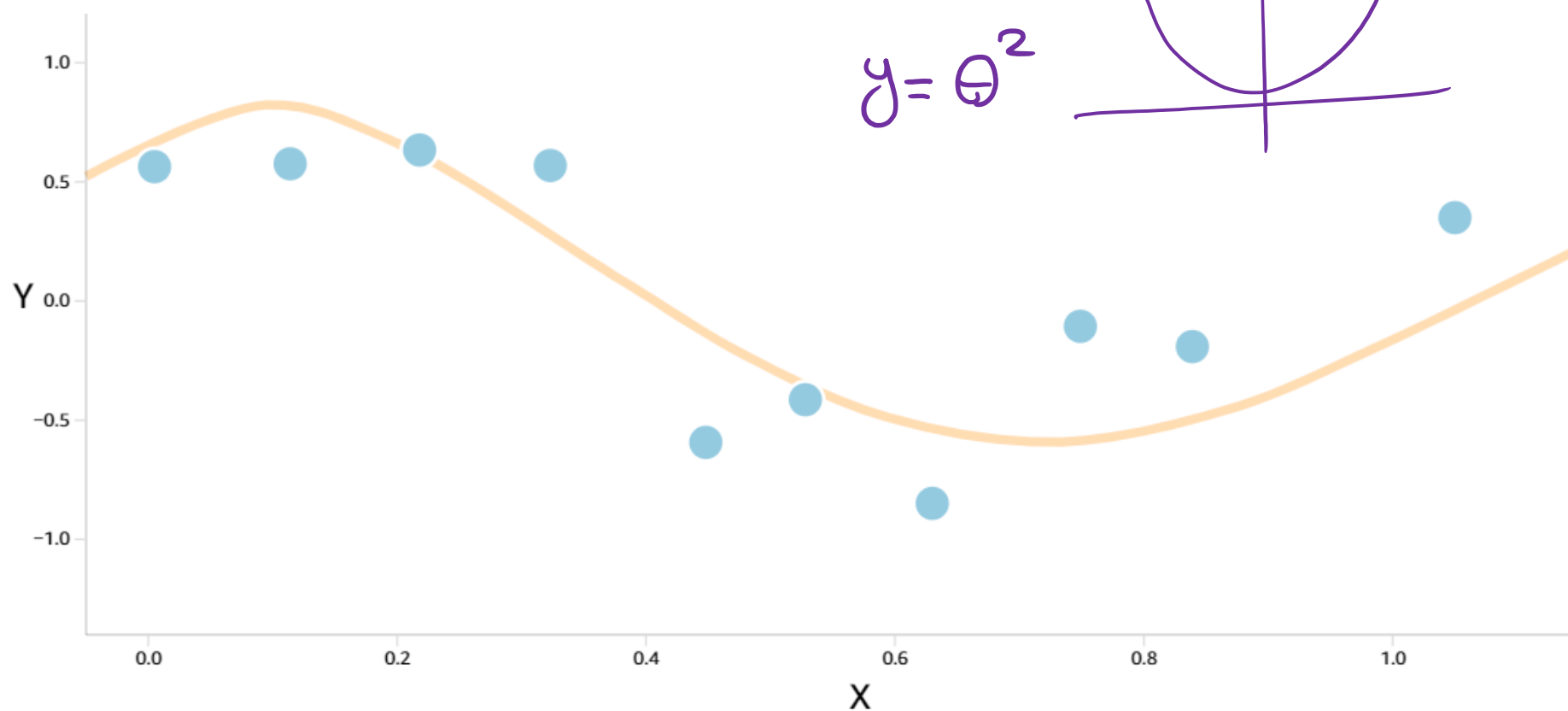
$$|\theta| \leq 100 = C$$

$$|\theta_0| + |\theta_1| + \dots + |\theta_d| \leq C$$

$$g(\theta) = |\theta|$$


$$\theta_0^2 + \theta_1^2 + \dots + \theta_d^2 \leq C$$

$$y = \theta^2$$


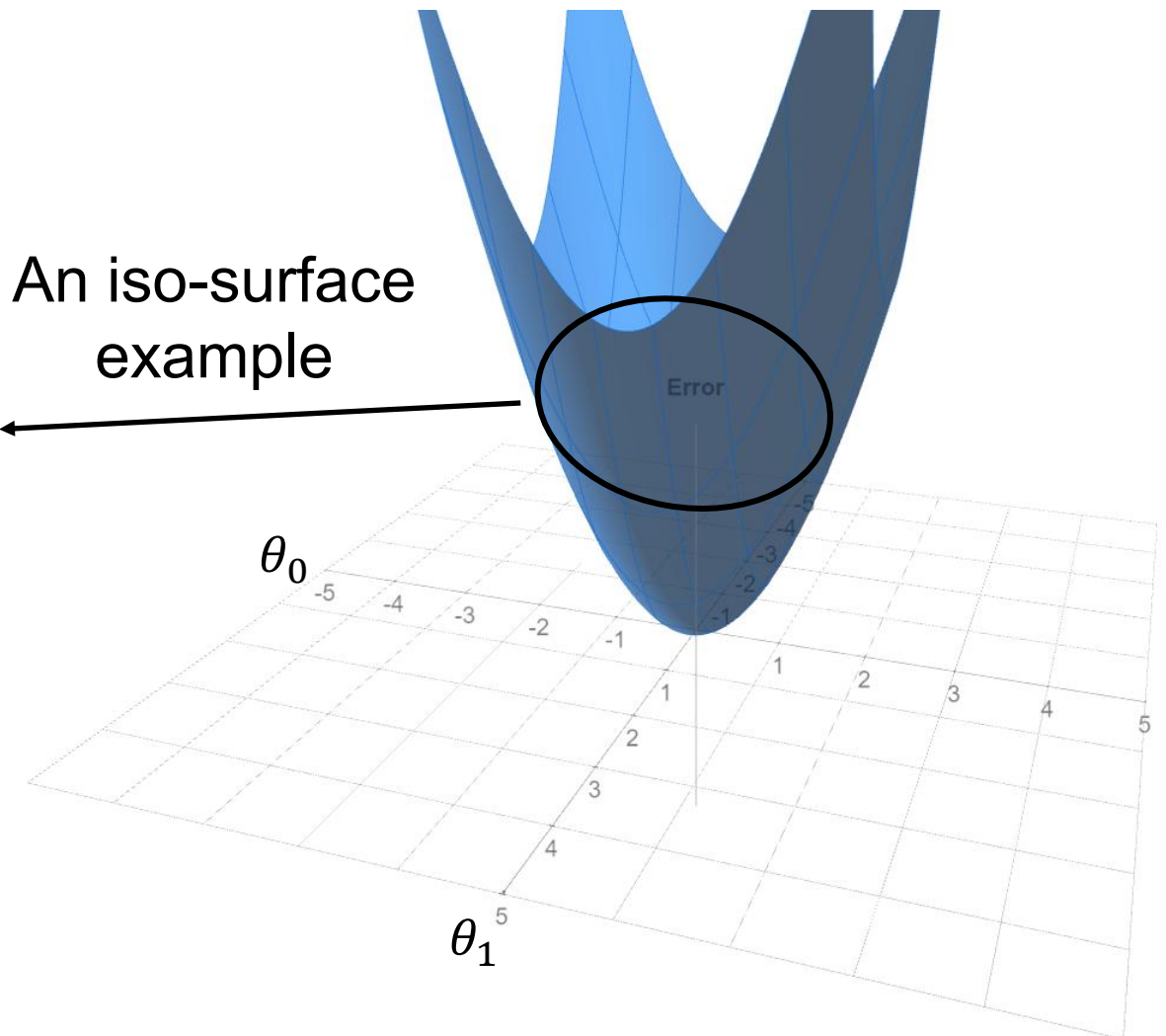


How can we get an optimal solution with a positive error for a model that overfits?

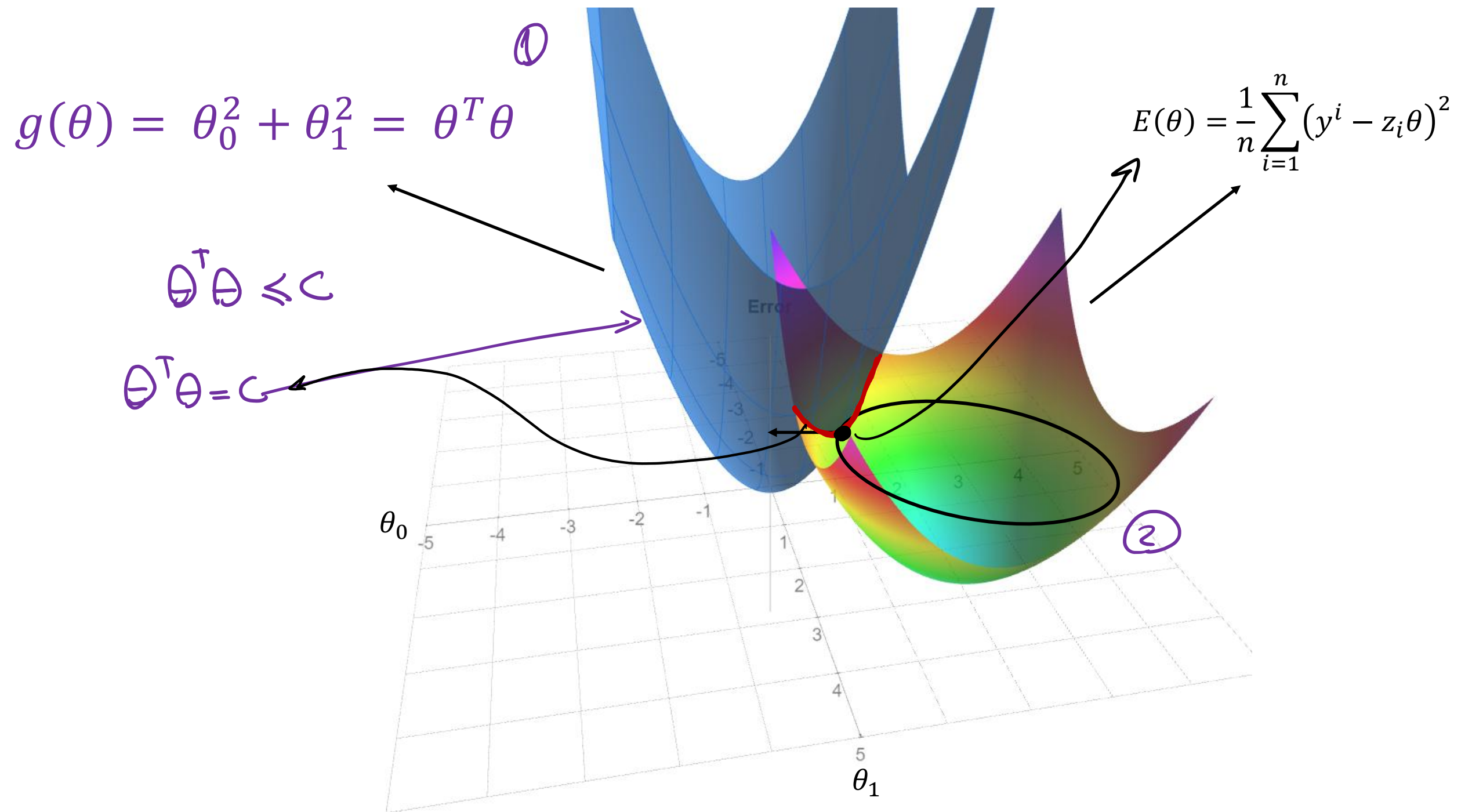
We need to introduce a constraint

$$g(\theta) = \theta_0^2 + \theta_1^2 = \theta^T \theta = C$$

An iso-surface
example



Error function together with a new introduced constraint



Let's define the Lagrange function

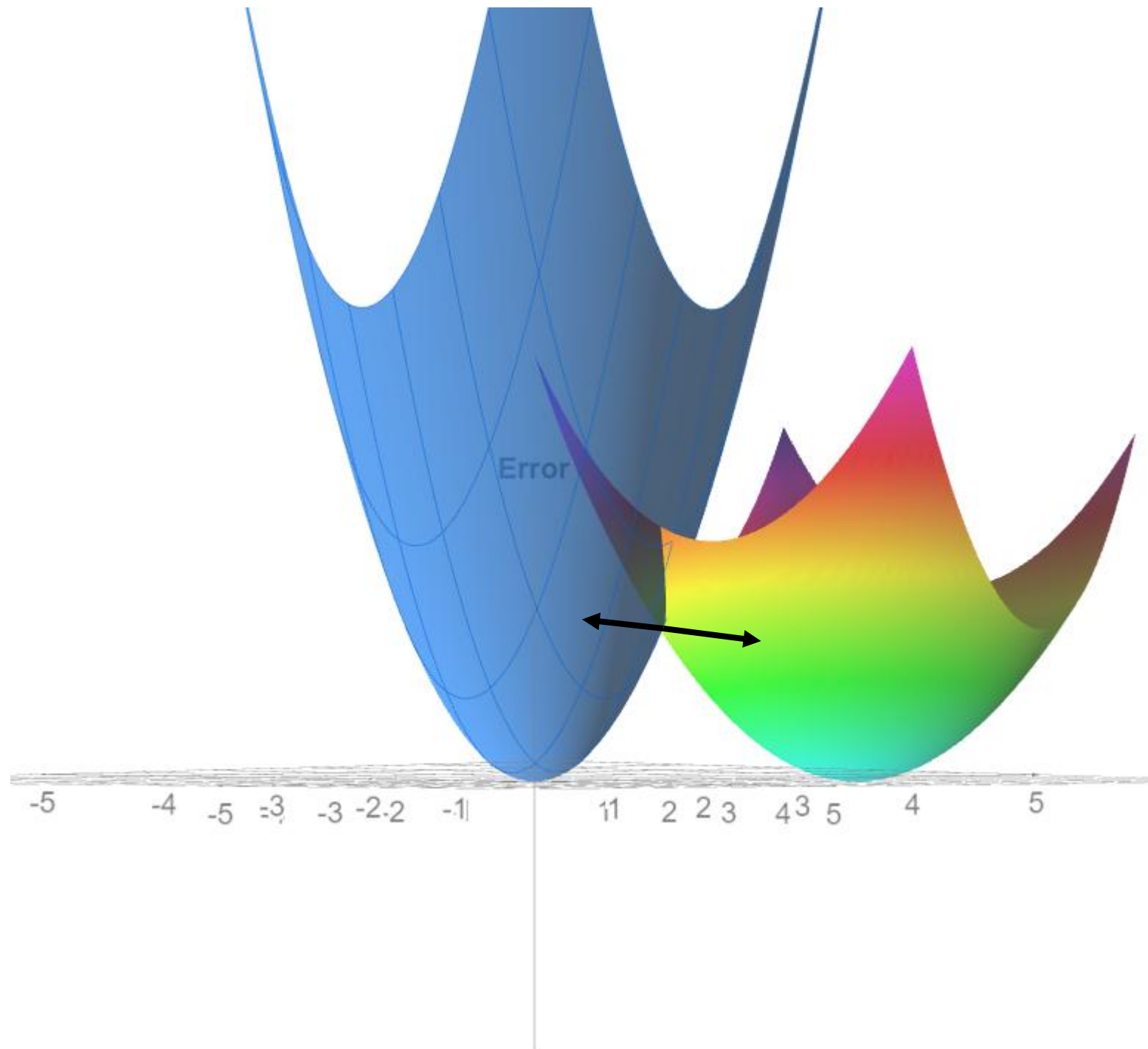
$$L(\theta, \lambda) = E(\theta) + \lambda g(\theta)$$

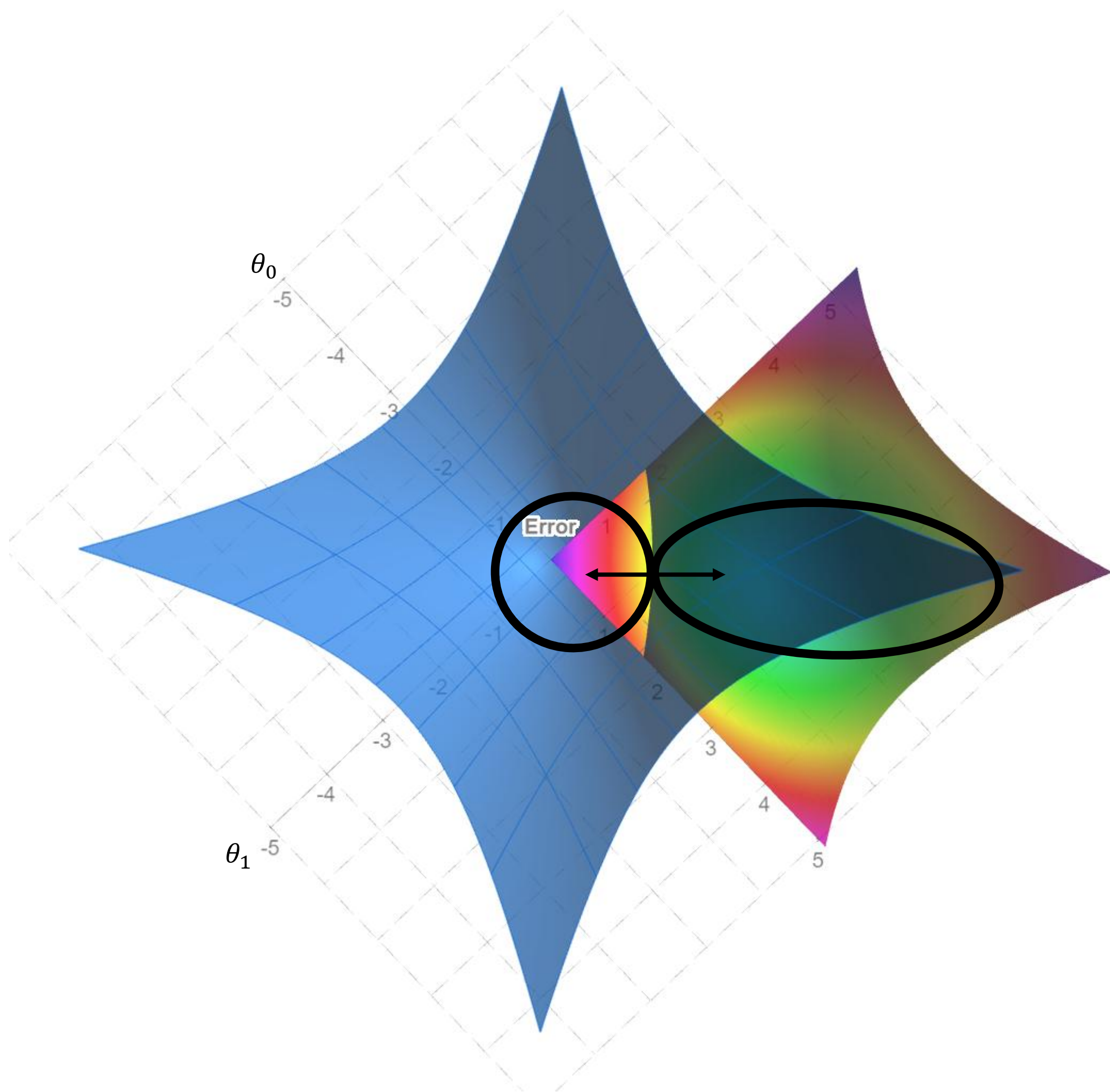
$$L(\theta, \lambda) = E(\theta) + \lambda \theta^T \theta$$

$$\nabla L(\theta, \lambda) = 0 \qquad \nabla [E(\theta) + \lambda \theta^T \theta] = 0$$

$$\nabla [E(\theta)] + \lambda \nabla [\theta^T \theta] = 0$$

How to enforce the gradient of Lagrange function to be zero





Let's calculate the gradients

Gradient of constraint $g(\theta)$

$$\nabla[\theta^T \theta] = 2\theta$$

$$\nabla[E(\theta)] + \lambda \nabla[\theta^T \theta] = 0$$

$$\nabla[E(\theta)] = -\lambda \nabla[\theta^T \theta]$$

$$\nabla E(\theta) = -2\lambda\theta$$

$$\nabla E(\theta) + 2\lambda\theta = 0$$

Let's do integration

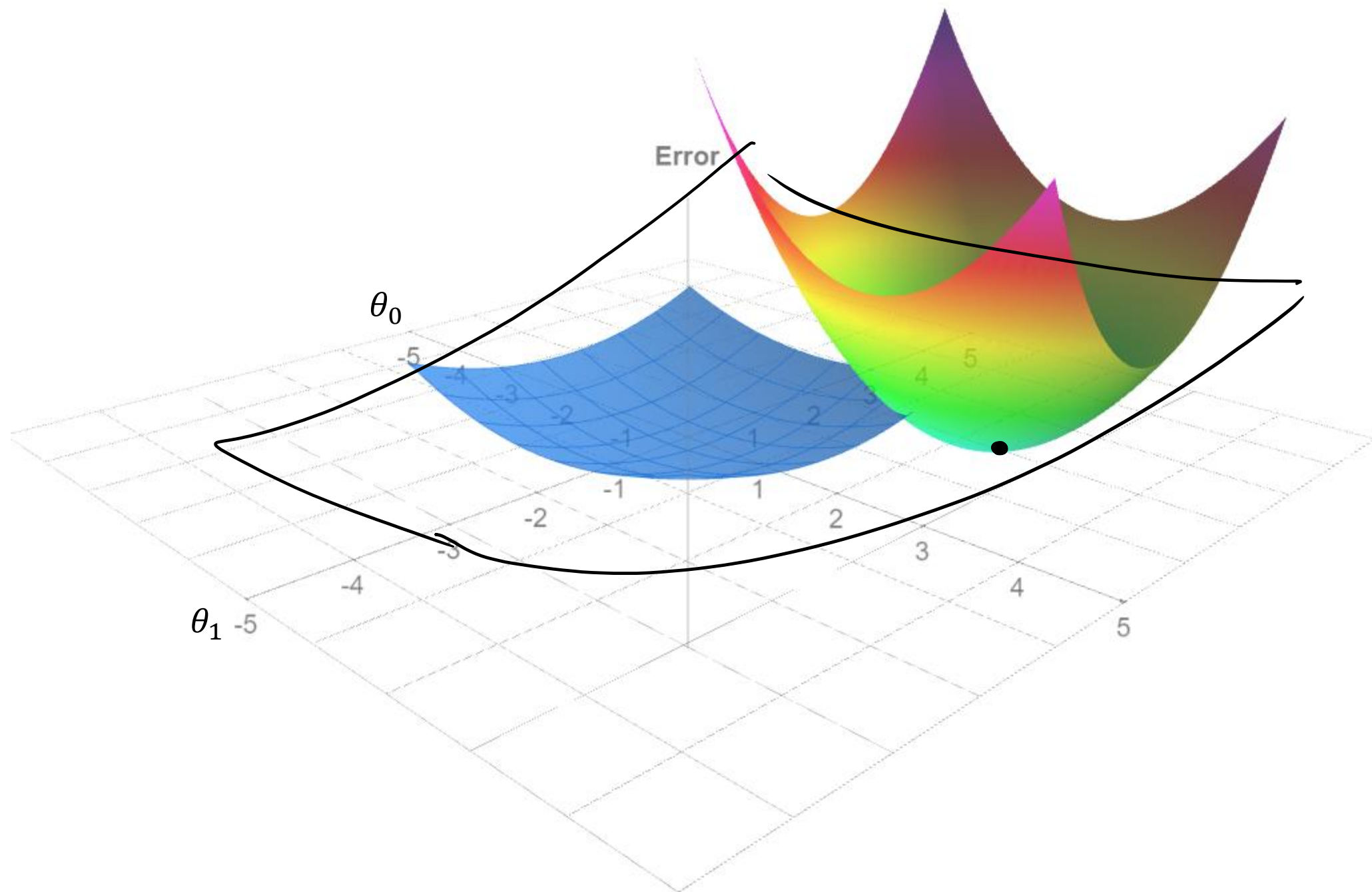
$$E(\theta) + \lambda\theta^T \theta$$

The effect of low Lambda

$$E(\theta) = \frac{1}{N} \sum (y_i - \hat{y}_i)^2$$

$$E(\theta) + \frac{\lambda}{N} \theta^T \theta$$

$C = \text{large}$



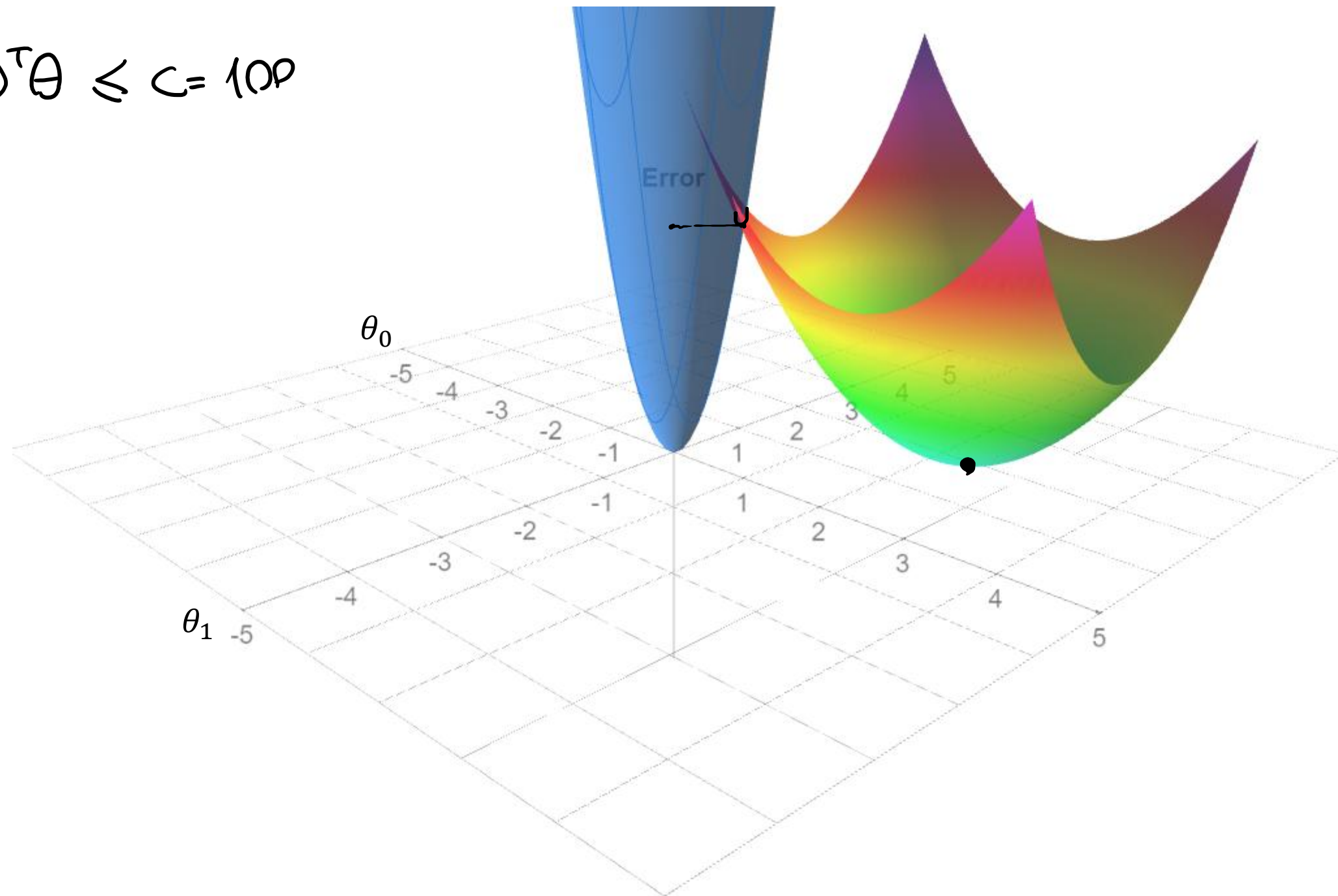
The effect of high Lambda

$$E(\theta) + \frac{\lambda}{N} \theta^T \theta$$

$C = \text{small}$

$$\theta^T \theta \leq 99$$

$$\theta^T \theta \leq C = 100$$



$$\text{Min } E(\theta) = \frac{1}{N} \sum_{i=1}^N (y_{a_i} - x_i^T \theta)^2$$

$$\|\theta\|^2 = \theta^T \theta \quad \|\theta\|^2 \leq C \rightarrow \text{s.t. } \theta^T \theta = C \quad g(\theta) = \theta^T \theta - C$$

$$L(\theta, \lambda) = E(\theta) + \lambda(\theta^T \theta - C) \Rightarrow \frac{\partial L}{\partial \theta} = 0 \quad ; \quad \frac{\partial L}{\partial \lambda} = 0$$

Let's say someone gave us λ already $\leadsto \lambda \rightarrow \text{constant}$

$$\text{Min } L(\theta) = E(\theta) + \lambda \theta^T \theta - \cancel{\lambda C} = E(\theta) + \lambda (\theta^T \theta)$$

$$\theta^T \theta = C$$

$$\theta_1^2 + \theta_2^2 + \dots + \theta_d^2 = C$$

$$\begin{array}{cc} \lambda \uparrow & \downarrow C = \theta^T \theta \\ \text{large} & \\ \downarrow \lambda & \theta^T \theta \\ \text{small} & \end{array}$$

Regularized Learning

Minimize $E(\theta) + \lambda \theta^T \theta$

overfit

Now we know Why this term leads to the regularization of parameters

Regularized Error

$$\tilde{E}(\theta) = \frac{1}{N} \sum_{i=1}^n (y^{\{i\}} - z^{\{i\}} \theta)^2 + \frac{\lambda}{2N} \|\theta\|_2^2$$

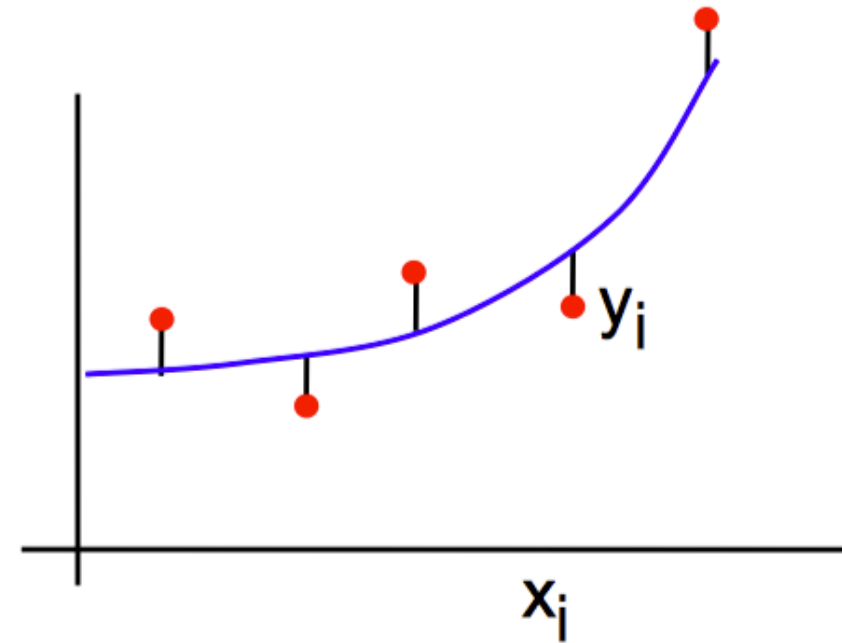
L2 Regularization term

Outline

- Overfitting and regularized learning
- Ridge regression 
- Lasso regression
- Determining regularization strength

Ridge Regression

$$\begin{aligned}\tilde{E}(\theta) \\ &= \frac{1}{N} \sum_{i=1}^n (y^{\{i\}} - z^{\{i\}} \theta)^2 + \lambda \|\theta\|_2^2\end{aligned}$$



$$\theta_0 + \theta_1 z_1 + \theta_2 z_2 + \cdots + \theta_d z_d + \epsilon = \mathbf{z} \boldsymbol{\theta}$$

General form

$$\tilde{E}(\theta) = \frac{1}{N} \sum_{i=1}^n (y^{\{i\}} - z^{\{i\}} \theta)^2 + \lambda \|\theta\|_2^2$$

$\sum \theta^T I \theta$
 $\sum |\theta|$

$$\theta^T D \theta$$

$$D = \begin{bmatrix} 0 & 0 & 0 & \dots \\ 0 & 1 & 0 & \dots \\ 0 & 0 & 1 & \dots \\ \vdots & \vdots & \vdots & \ddots \end{bmatrix}$$

Matrix form

$$\tilde{E}(\theta) = \frac{1}{N} (y - z\theta)^T (y - z\theta) + \frac{\lambda}{2} \|\theta\|_2^2$$

$$\theta_1 x + \theta_2 x^2 + (\theta_3) x^3$$

$$\frac{\partial \tilde{E}(\theta)}{\partial \theta} = -z^T (y - z\theta) + \lambda \theta$$

$$\frac{\partial (\frac{\lambda}{2} \theta^2)}{\partial \theta} = \lambda \theta$$

$$(z^T z + \lambda I) \theta = z^T y$$

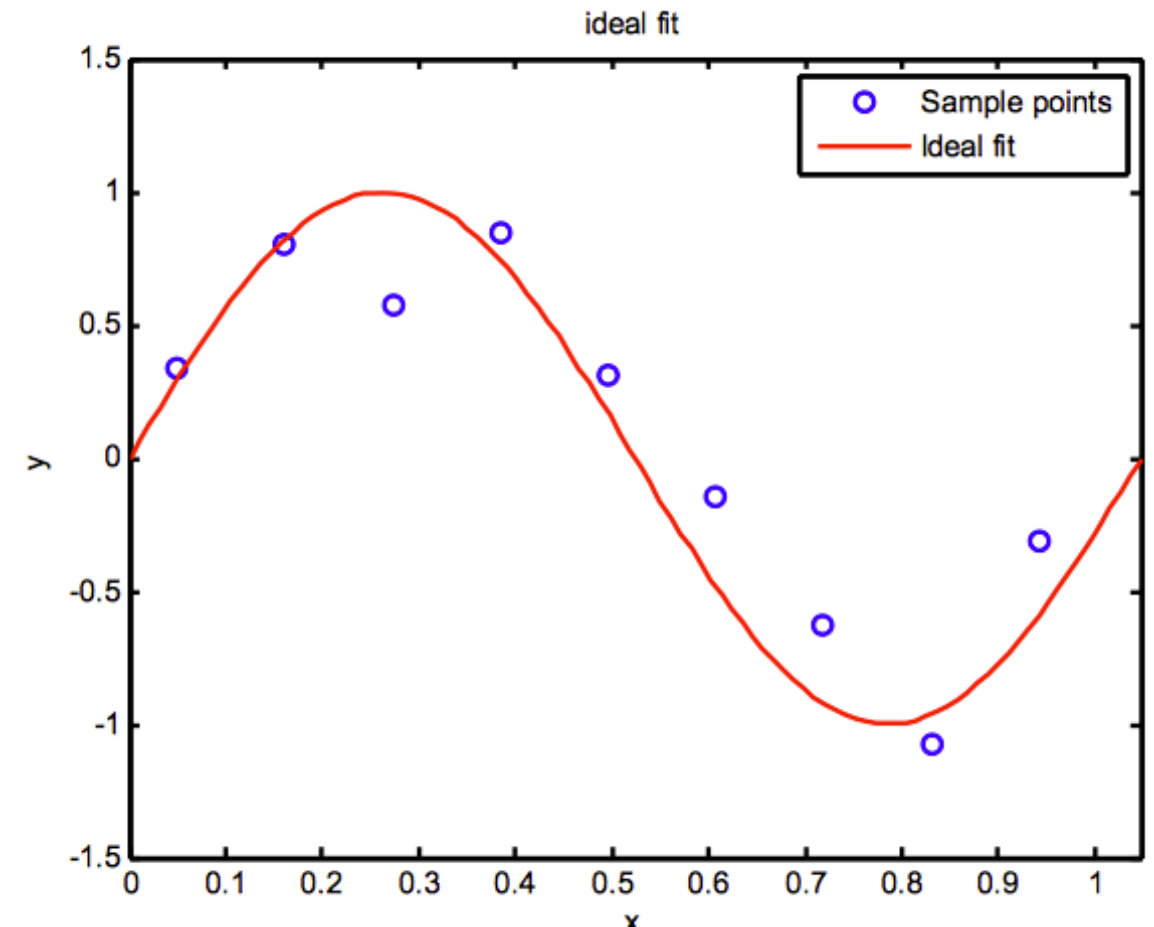
$$\rightarrow \theta = (z^T z + \lambda I)^{-1} z^T y$$

$$\theta = \underset{\text{overfitted}}{(z^T z)^{-1}} z^T y$$

$$\frac{1}{z^T z + \lambda I}$$

Ridge Regression Example

- The red curve is the true function (which is not a polynomial)
- The data points are samples from the curve with added noise in y.
- There is a choice in both the degree, D , of the basis functions used, and in the strength of the regularization

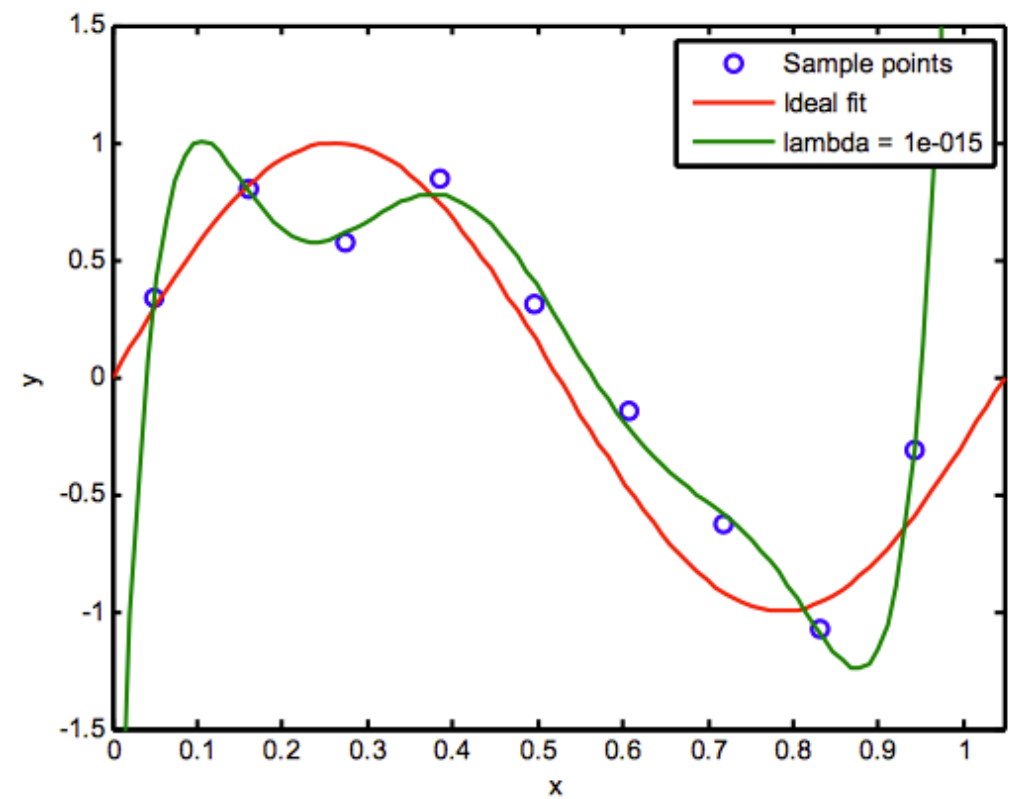
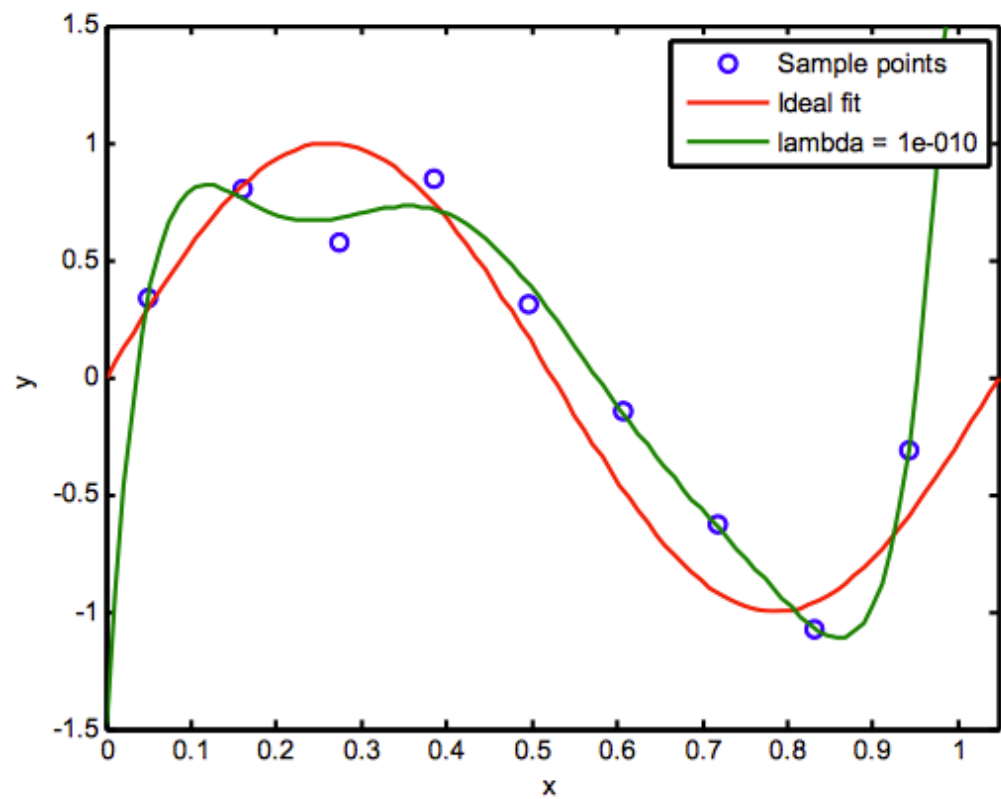
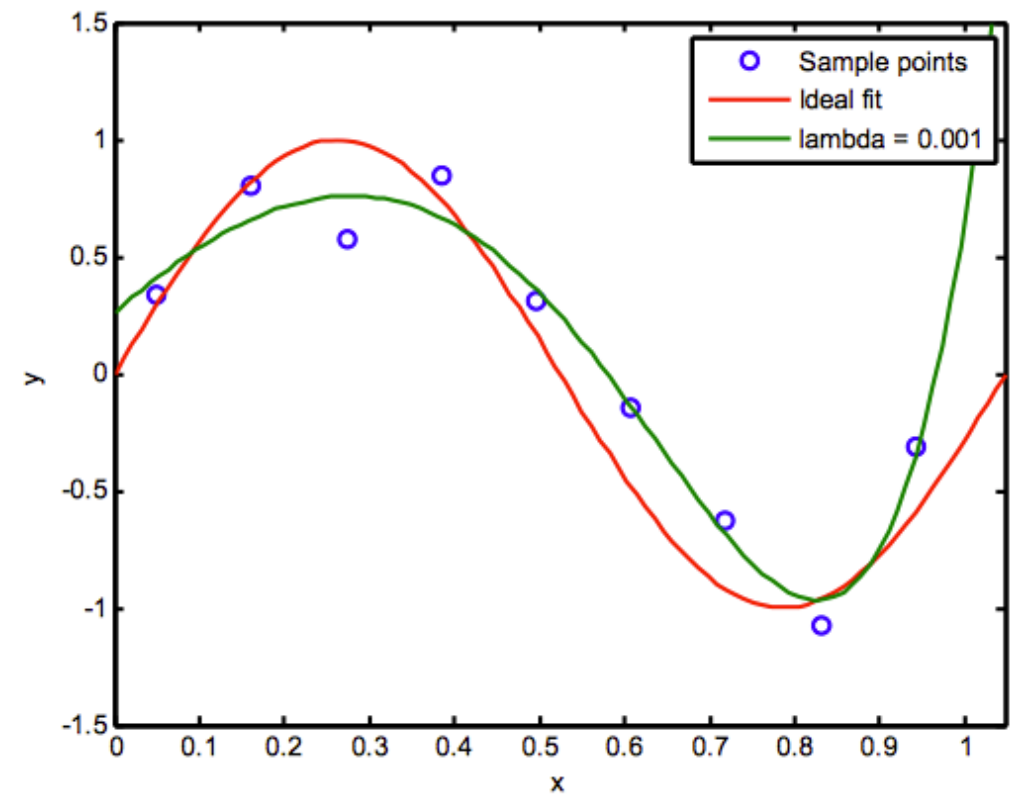
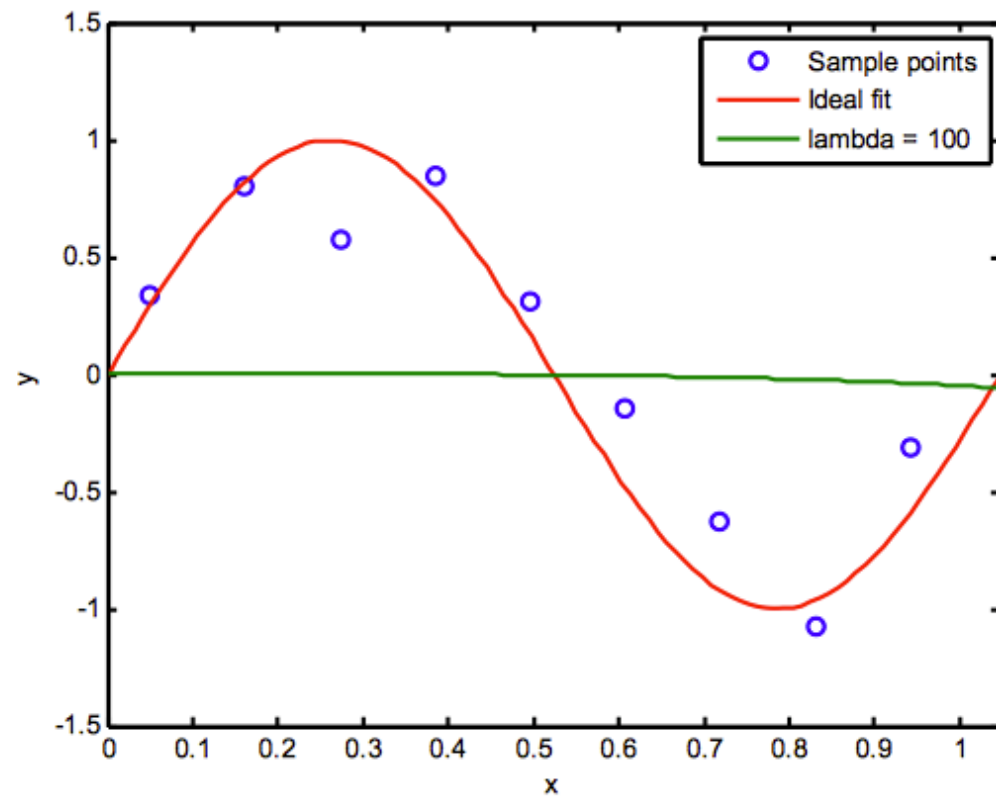


$$f(x, \theta) = z\theta$$

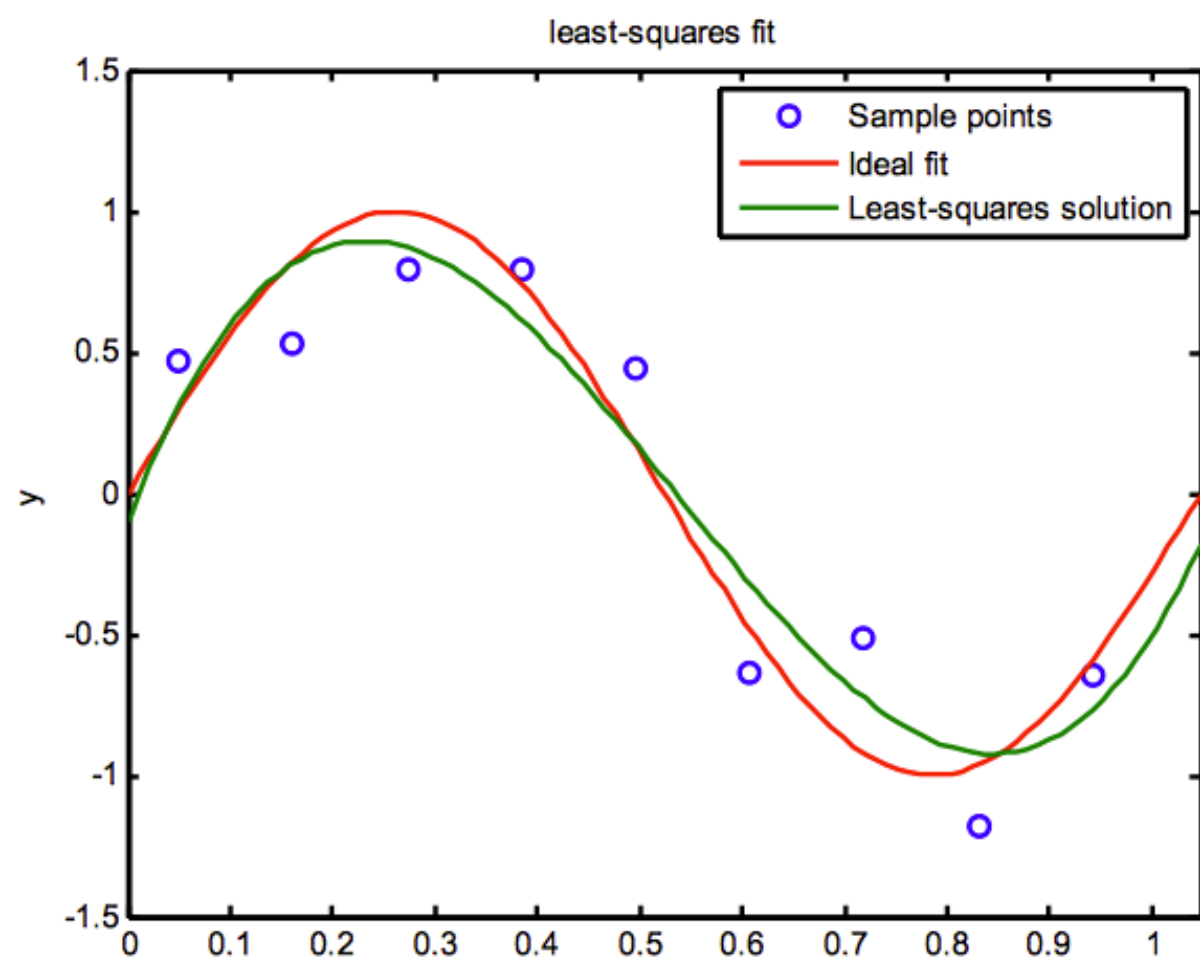
$$z: x \rightarrow z$$

$$\tilde{E}(\theta) = \frac{1}{N} \sum_{i=1}^n (y^{\{i\}} - z^{\{i\}}\theta)^2 + \lambda \|\theta\|_2^2 \quad \theta \in \mathbb{R}^{D+1}$$

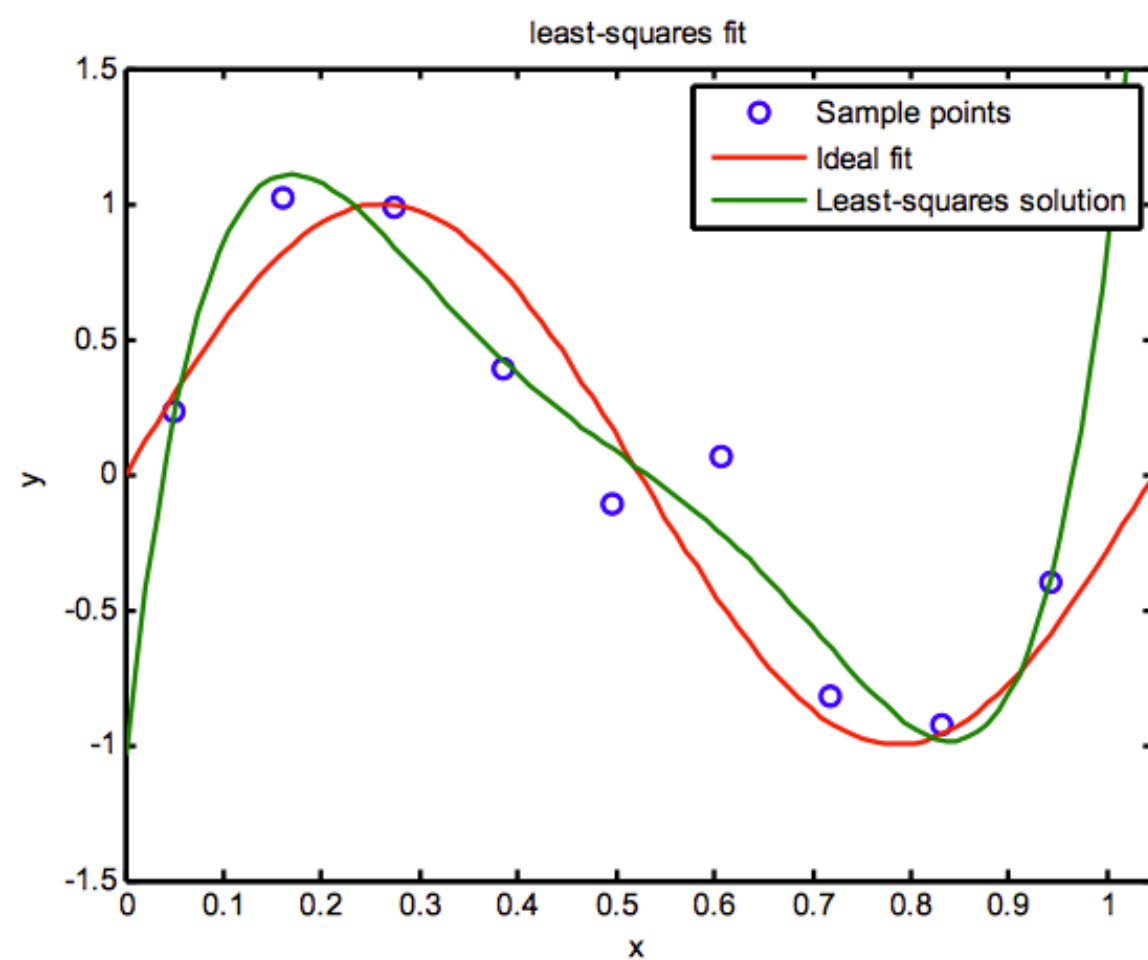
$N = 9$ samples, $D = 7$



$D = 3$



$D = 5$



Outline

- Overfitting and regularized learning
- Ridge regression
- Lasso regression 
- Determining regularization strength

Regularized Regression

$$\tilde{E}(\theta) = \frac{1}{N} \sum_{i=1}^n (y^{\{i\}} - z^{\{i\}} \theta)^2 + \lambda \|\theta\|_2^2$$

Squared loss/Error

$$\frac{1}{N} \sum_{i=1}^n (y^{\{i\}} - z^{\{i\}} \theta)^2$$

L2 Regularizer

$$\lambda \|\theta\|_2^2$$

Now let's look at another regularization choice.

The Lasso Regularization (L1 norm) and sparsity

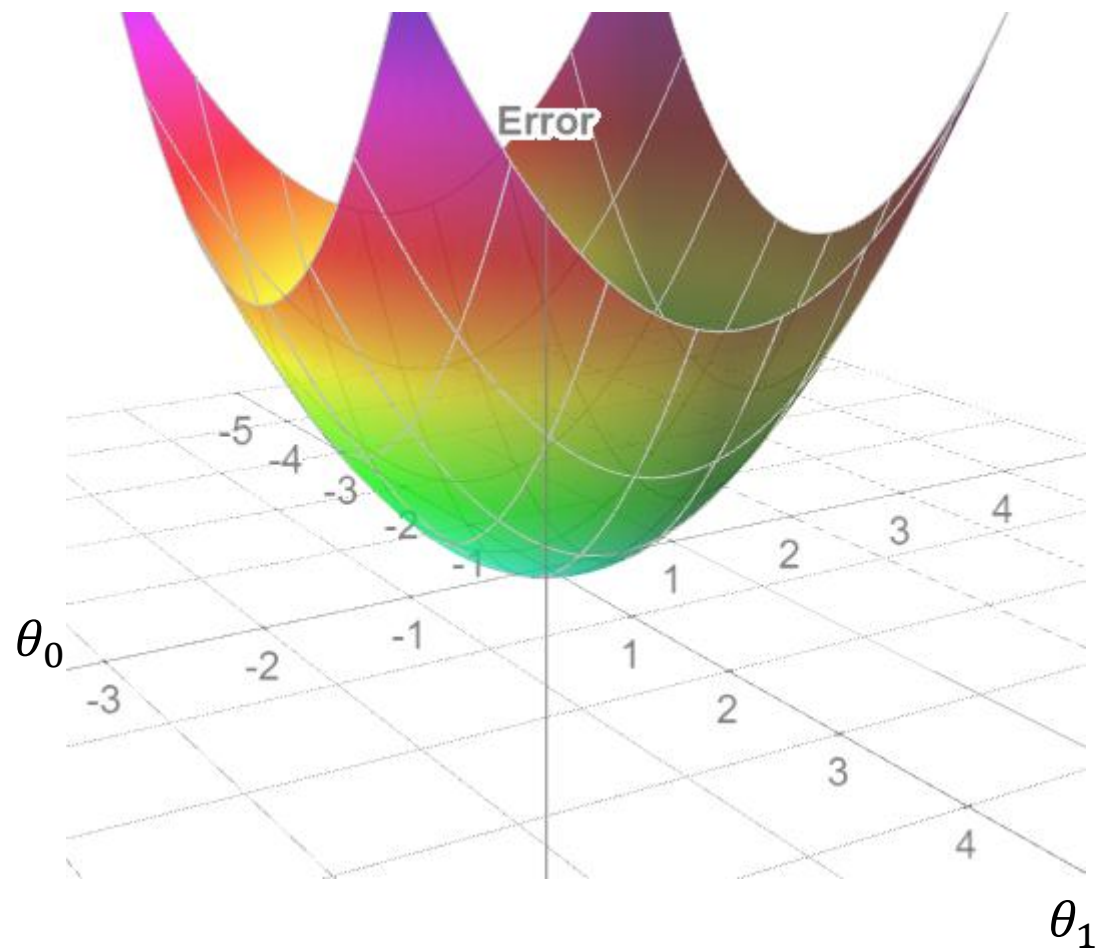
Lasso = Least Absolute Shrinkage and Selection Operator

$$\tilde{E}(\theta) = \frac{1}{N} \sum_{i=1}^n (y^{\{i\}} - z^{\{i\}} \theta)^2 + \lambda \|\theta\|_1$$

L1 norm induces sparsity. This means that some of the weights become zero, and the feature contribution will be completely removed. L1 Regularizer could be used for feature selection

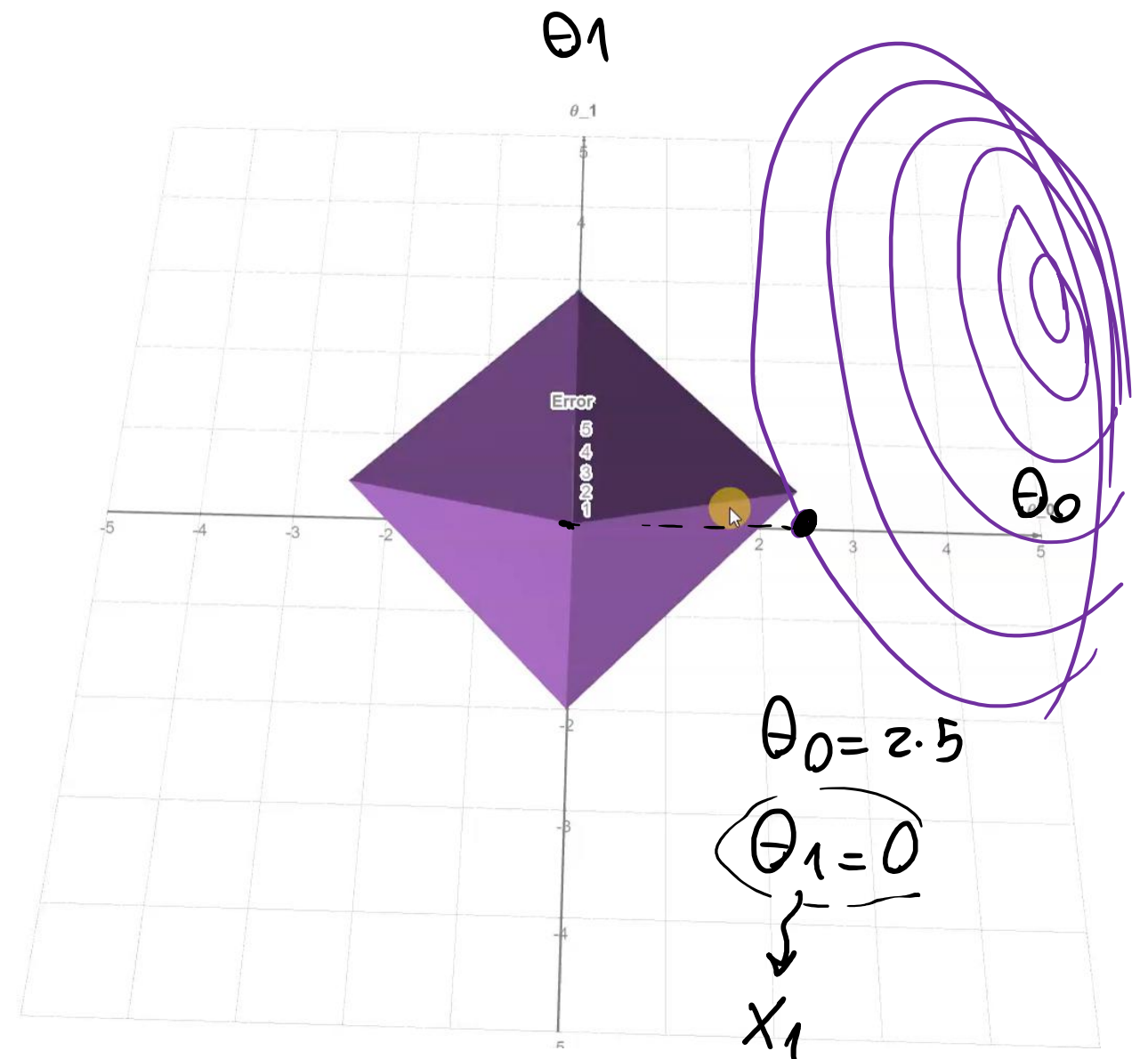
Ridge Regularizer

$$g(\theta) = \theta_0^2 + \theta_1^2 = \theta^T \theta$$



Lasso Regularizer

$$g(\theta) = \theta_0 + \theta_1 = \theta$$



[Animation](#)

Let's say we have two parameters (θ_0 and θ_1)

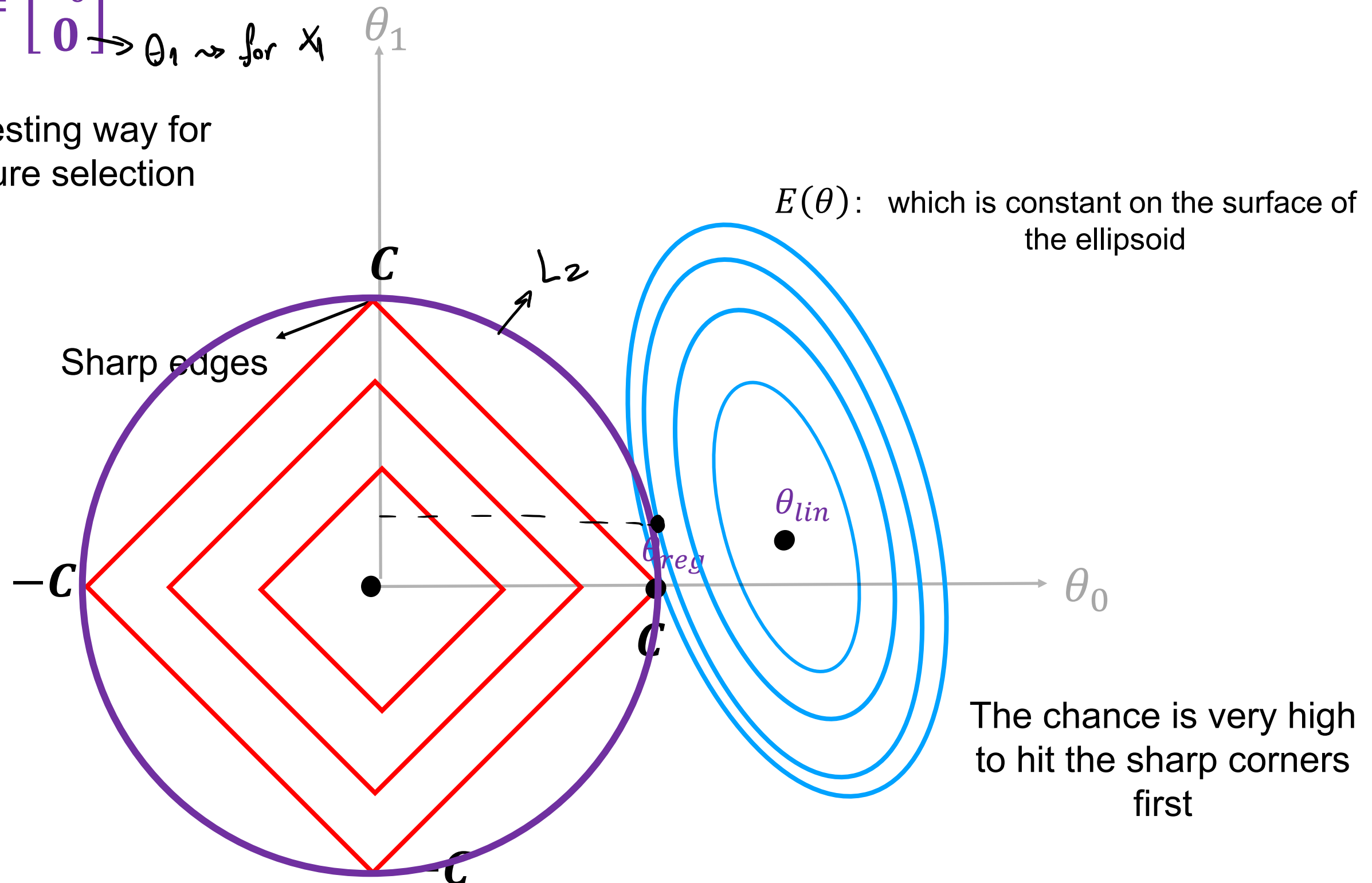
$$\text{Min } E(\theta) = \frac{1}{N} (z\theta - y)^T (z\theta - y) + \lambda \|\theta\|_1$$

$$\hat{\theta} = \begin{bmatrix} \theta_0 \\ \theta_1 \end{bmatrix}$$

$\theta_1 \rightsquigarrow$ for x_1

biase term

Interesting way for
feature selection



Ridge versus Lasso

Ridge

$$\tilde{E}(\theta) = \frac{1}{N} (y - z\theta)^T (y - z\theta) + \lambda \|\theta\|_2^2$$

It is a convex model

Both mean squared error
and L2 regularizer are
differentiable.

We can get a closed form
solution

Lasso



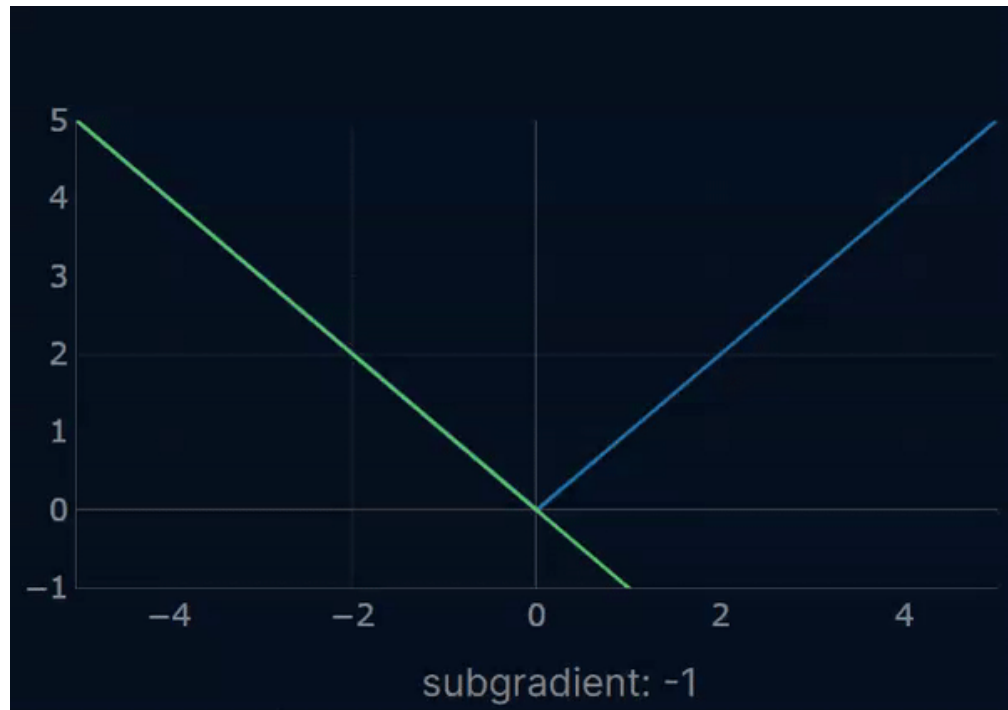
$$\tilde{E}(\theta) = \frac{1}{N} (y - z\theta)^T (y - z\theta) + \lambda \|\theta\|_1$$

It is a convex model

L1 regularizer is NOT
differentiable.

We can **NOT** get a closed
form solution

Sub-gradient Descend in Lasso



$$\theta^{t+1} \leftarrow \theta^t - \alpha \frac{\partial \tilde{E}(\theta)}{\partial \theta}$$

$$\theta = \begin{bmatrix} 0.07 \\ \vdots \end{bmatrix}$$

$$\tilde{E}(\theta) = \frac{1}{N} (y - z\theta)^T (y - z\theta) + \lambda \|\theta\|_1$$

$$\frac{\partial \tilde{E}(\theta)}{\partial \theta} = -z^T (y - z\theta) + \frac{\partial (\lambda \|\theta\|_1)}{\partial \theta}$$

Using Sub-gradient

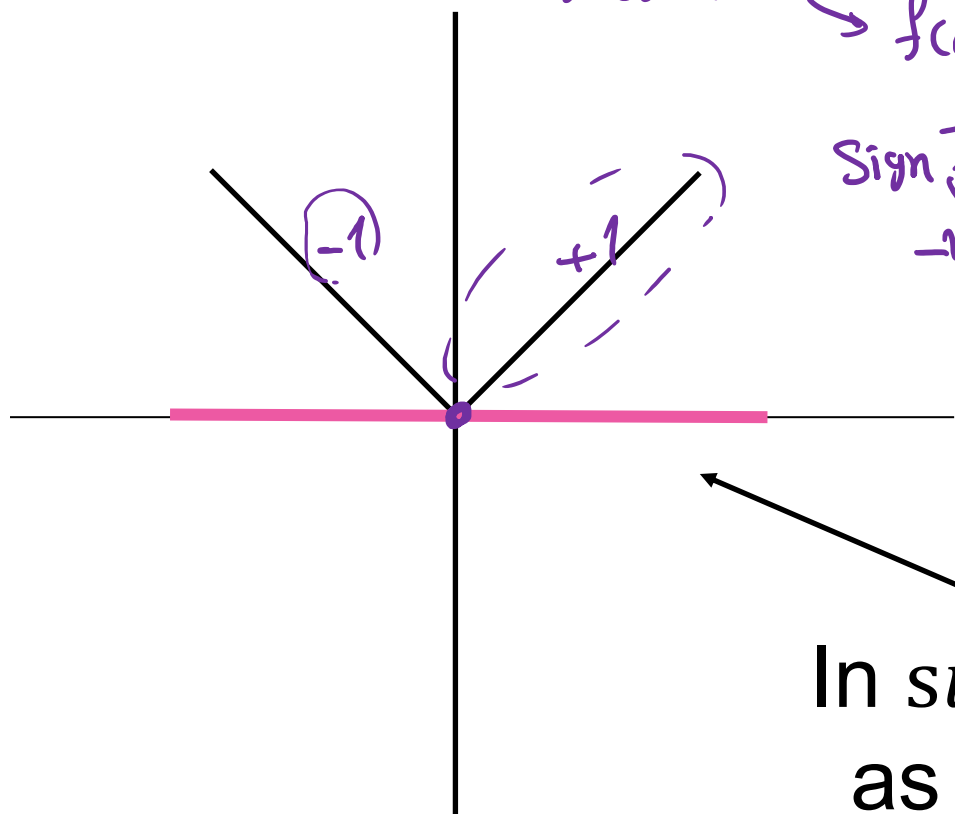
$$\frac{\partial \tilde{E}(\theta)}{\partial \theta} = -z^T (y - z\theta) + \lambda \text{sign}(\theta) = 0$$

$$\theta = ?$$

$$f(\theta) = |\theta| \begin{cases} \rightarrow f(\theta) = \theta \rightarrow f'(\theta) = 1 \\ \rightarrow f(\theta) = -\theta \rightarrow f'(\theta) = -1 \end{cases}$$

$$\text{Sign} \begin{cases} + \rightarrow +1 \\ 0 \rightarrow 0 \\ - \rightarrow -1 \end{cases}$$

$$\frac{\partial |\theta|}{\partial \theta} = \text{sign}(\theta)$$



In *sign* function, we use this sub-gradient line as our under-estimator (below our function)

Outline

- Overfitting and regularized learning
- Ridge regression
- Lasso regression
- Determining regularization strength 

we need optimize hyper parameters

Training
%70

Validation
%10

Testing
%20

Leave-One-Out Cross Validation

$$S = 0.01$$

$$S = 0.03 \quad Z \quad Y$$

$$S = 0.02$$

For every $i = 1, \dots, n$: $\theta = \left(Z^T Z + S I \right)^{-1} Z^T Y$

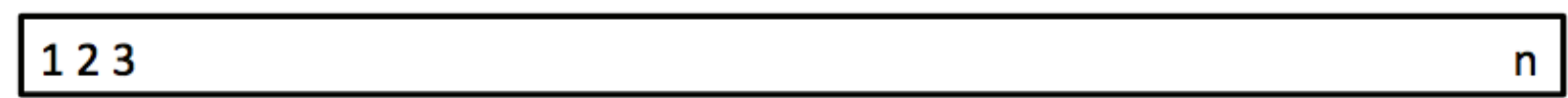
$\begin{matrix} d+1 \times 1 & & d+1 \times d+1 & & d+1 \times n & & n \times 1 \end{matrix}$

$$E(\theta) = \frac{1}{N} \sum_{i=1}^N (y_i - x_i^T \theta)^2$$

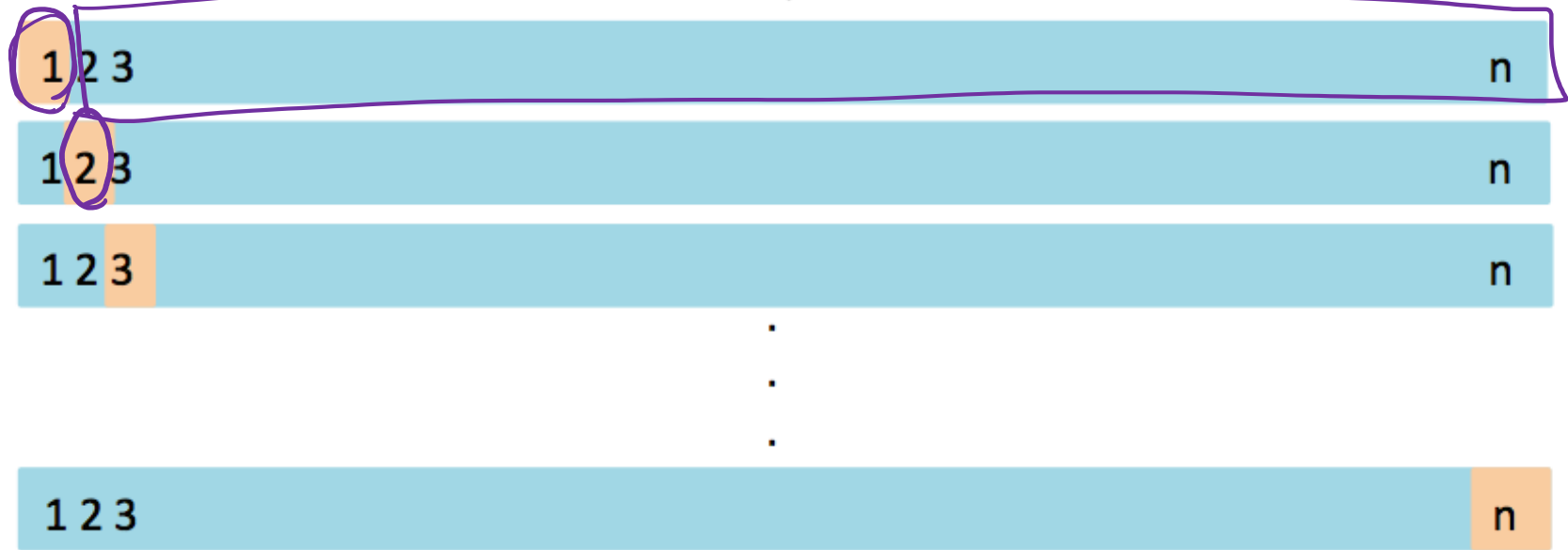
- ▶ train the model on every point except i ,
- ▶ compute the test error on the held out point.

Average the test errors.

$$CV_{(n)} = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i^{(-i)})^2$$



99 datapoints



avg error

$$\Rightarrow \frac{1}{100} (E(\theta_1) + \dots + E(\theta_n))$$

$$E(\theta_1) = (y_1 - x_1^T \theta)^2$$

$$E(\theta_2)$$

$$E(\theta_{100})$$

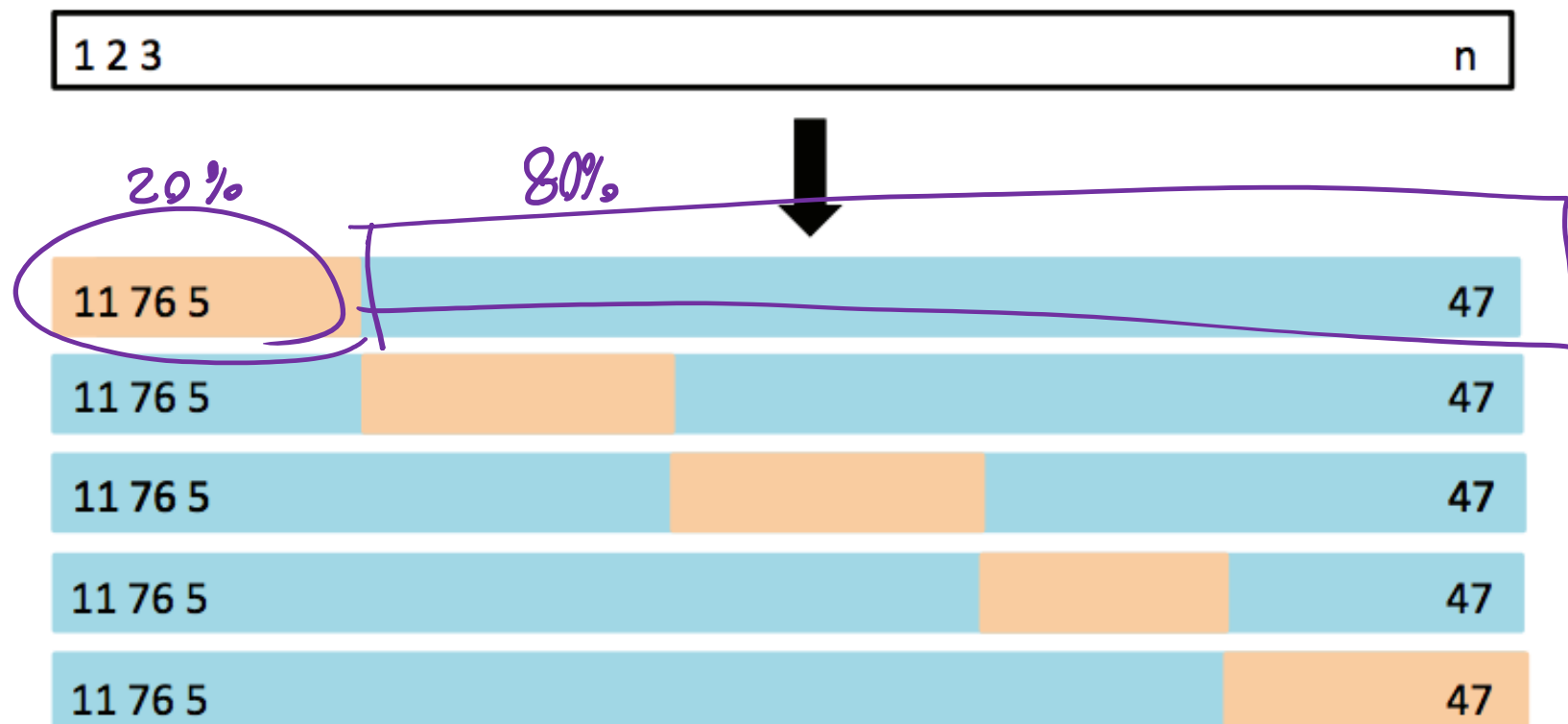
K-Fold Cross Validation

Split the data into k subsets or *folds*.

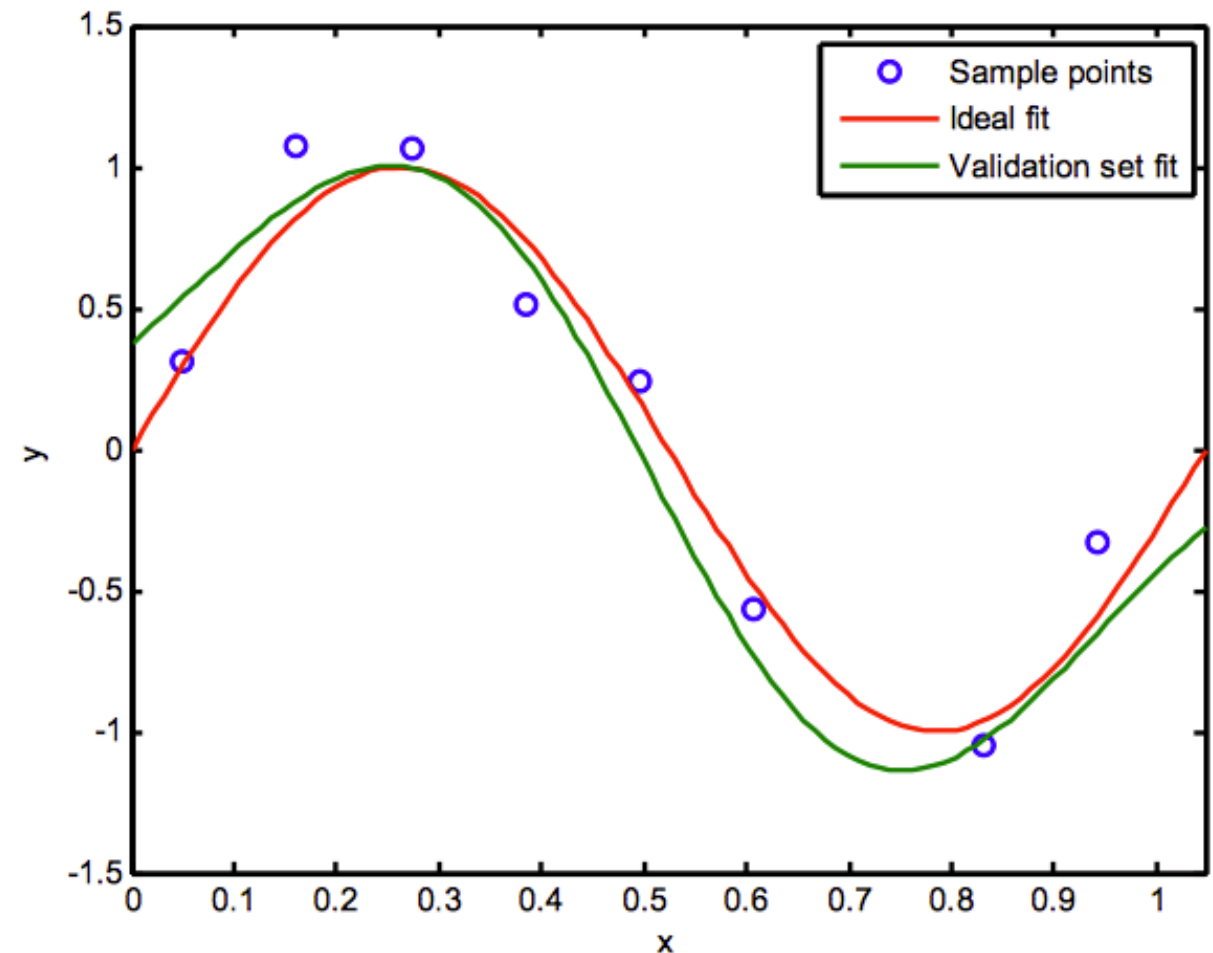
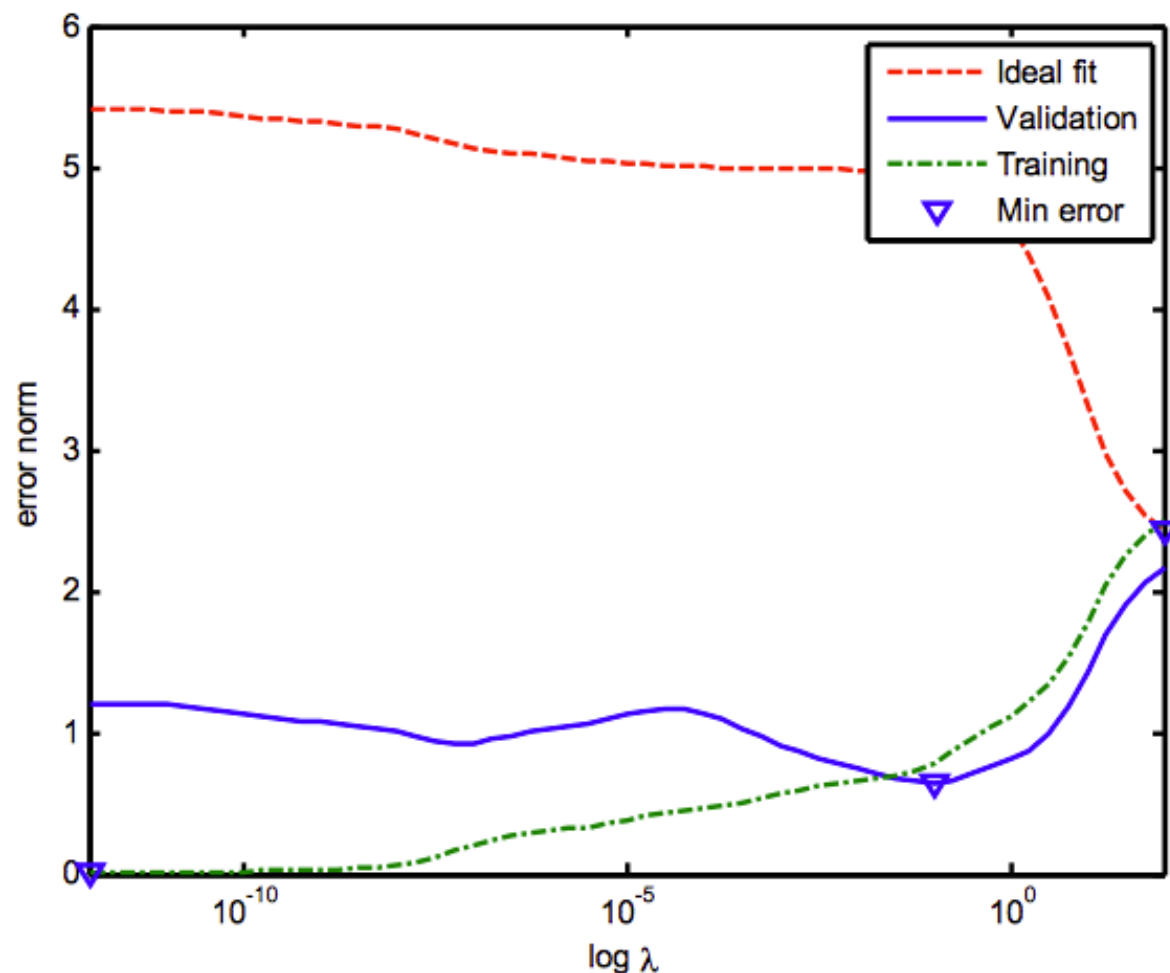
For every $i = 1, \dots, k$:

- ▶ train the model on every fold except the i th fold,
- ▶ compute the test error on the i th fold.

Average the test errors.



Choosing λ Using Validation Dataset



Pick up the lambda with the lowest mean value of rmse calculated by Cross Validation approach

Take-Home Messages

- What is overfitting
- What is regularization
- How does Ridge regression work
- Sparsity properties of Lasso regression
- How to choose the regularization coefficient λ