



Machine Learning: The Optimization Perspective

Konstantin Tretyakov

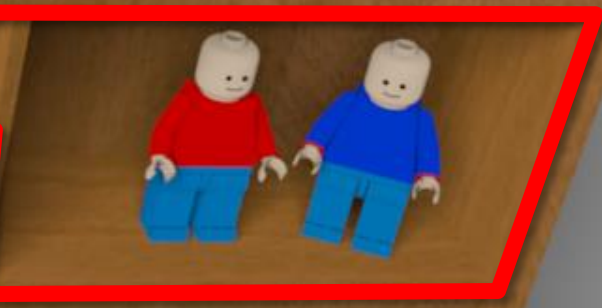
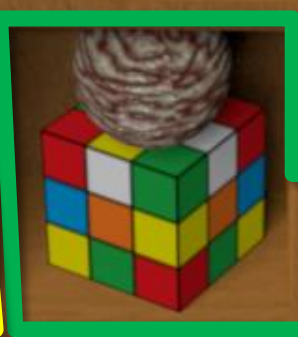
<http://kt.era.ee>

**IFI Summer
School 2014**

STACC

Software Technology and
Applications Competence Center





So far...

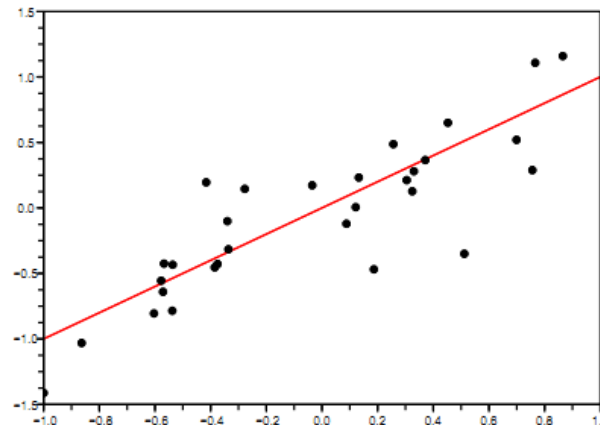
- ▶ Machine learning is important and interesting
- ▶ The general concept:

Fitting models to data

So far...

- ▶ Machine learning is important and interesting
- ▶ The general concept:

**Searching for the best
model fitting to data**

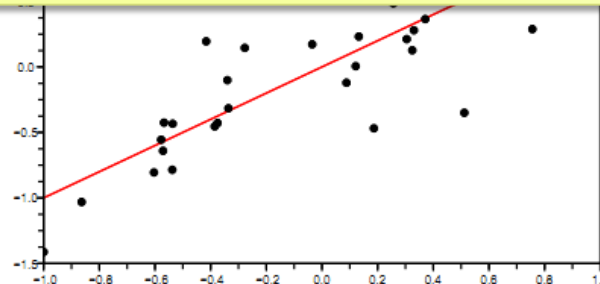


So far...

- ▶ Machine learning is important and interesting
- ▶ The g

Optimization

**Probability
Theory**



The Land of Machine Learning

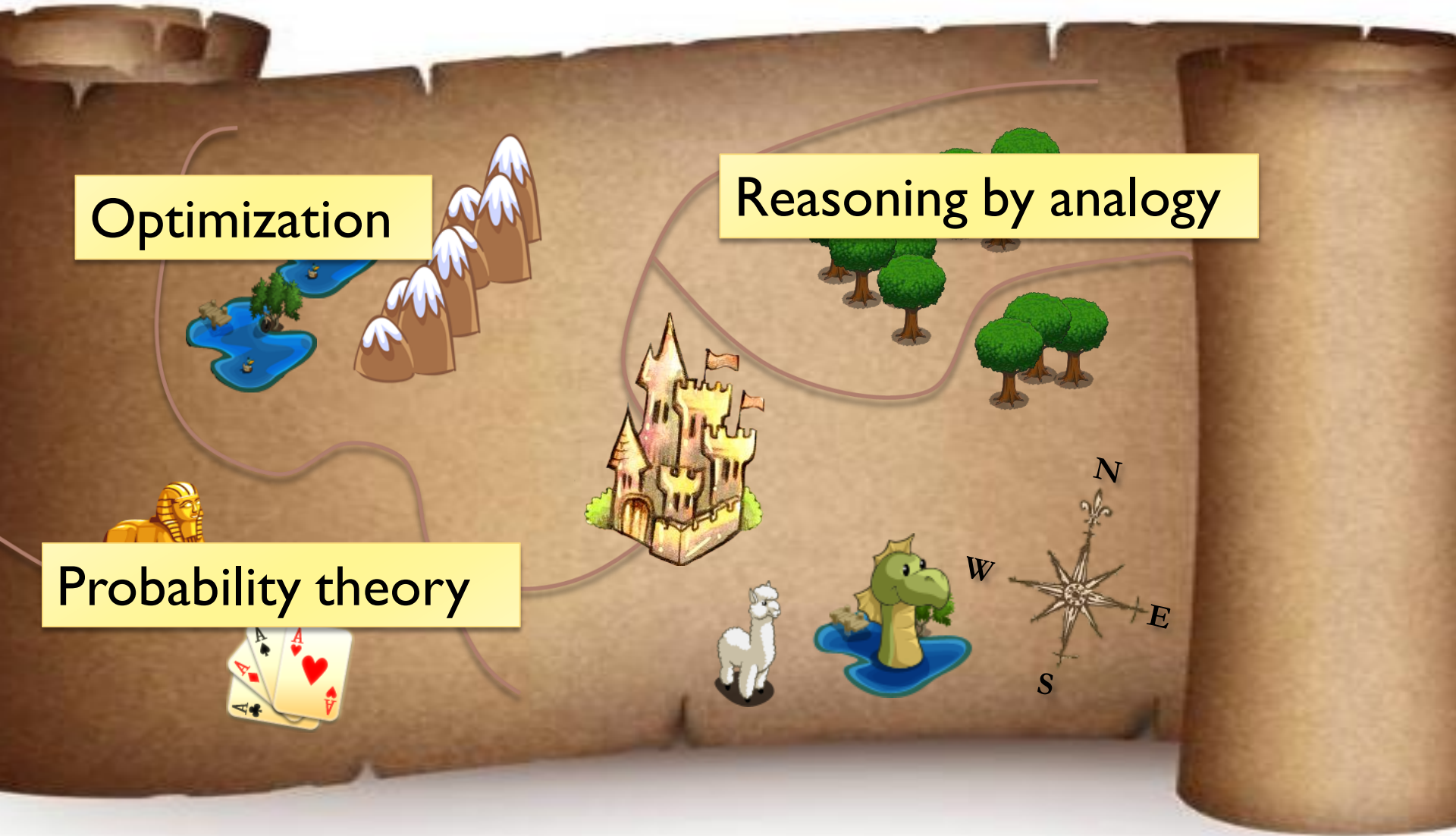


The Land of Machine Learning

Optimization

Reasoning by analogy

Probability theory



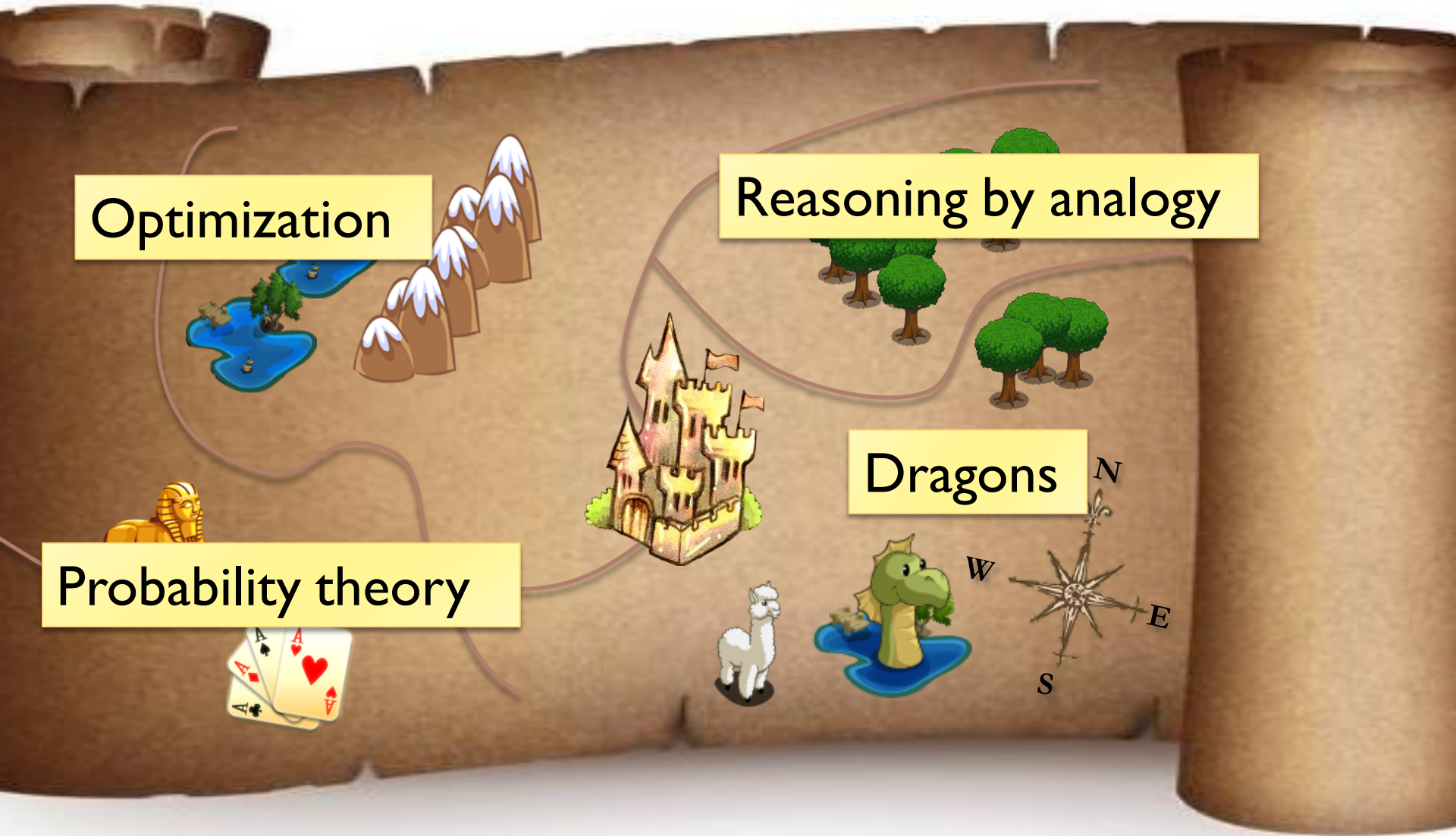
The Land of Machine Learning

Optimization

Reasoning by analogy

Probability theory

Dragons



What do you need to know about optimization?



What do you need to know about optimization?



1. Optimization is **important**
2. Optimization is **possible**



What do you need to know about optimization?



1. Optimization is **important**
2. Optimization is **possible***

* Basic **techniques**

- ▶ Constrained / Unconstrained
- ▶ Analytic / Iterative
- ▶ Continuous / Discrete



Special cases of optimization

- ▶ Machine learning

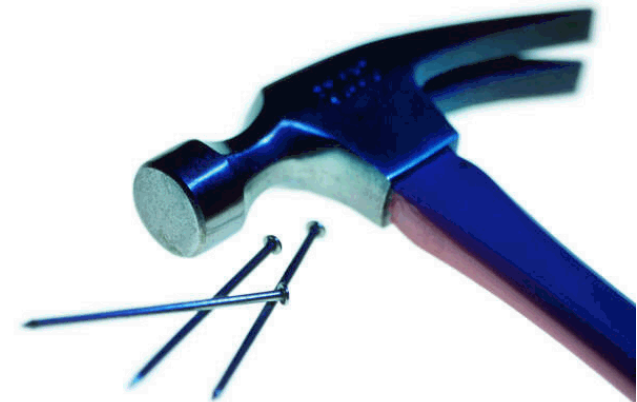
- ▶ ...

Special cases of optimization

- ▶ Machine learning
- ▶ Algorithms and data structures
- ▶ General problem-solving
- ▶ Management and decision-making

Special cases of optimization

- ▶ Machine learning
- ▶ Algorithms and data structures
- ▶ General problem-solving
- ▶ Management and decision-making
- ▶ Evolution
- ▶ *The Meaning of Life?*



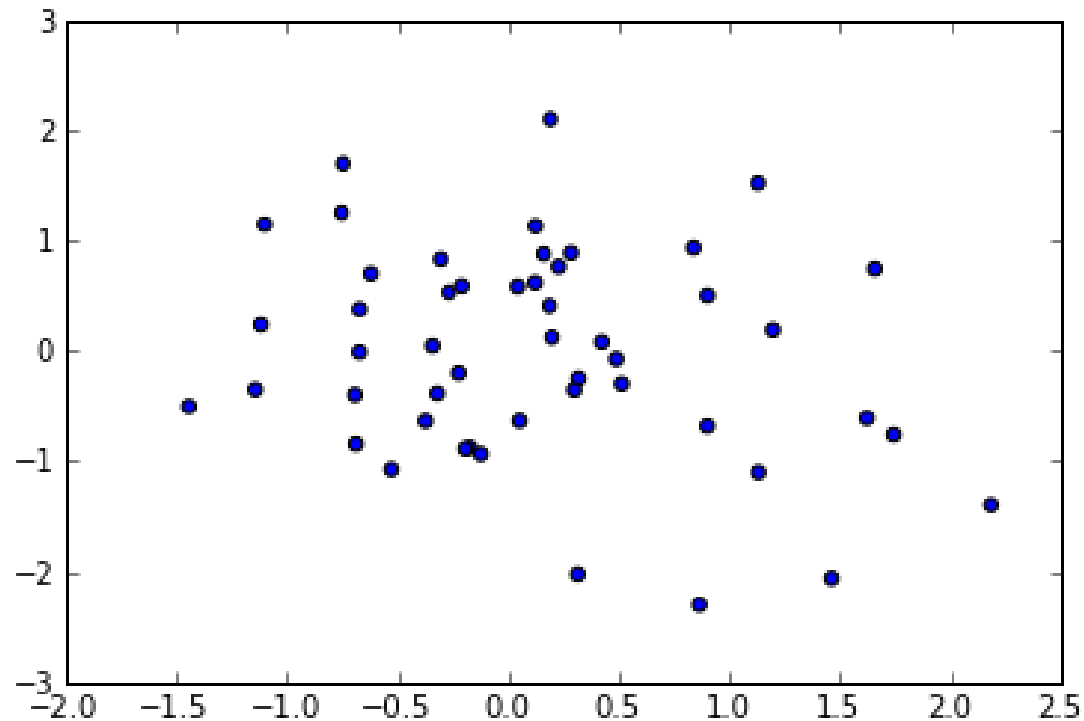
Example

Problem. Given a dataset $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in \mathbf{R}^m$, find $\mathbf{w}$, that minimizes

$$f(\mathbf{w}) = \sum_{i=1}^n \|\mathbf{x}_i - \mathbf{w}\|^2$$

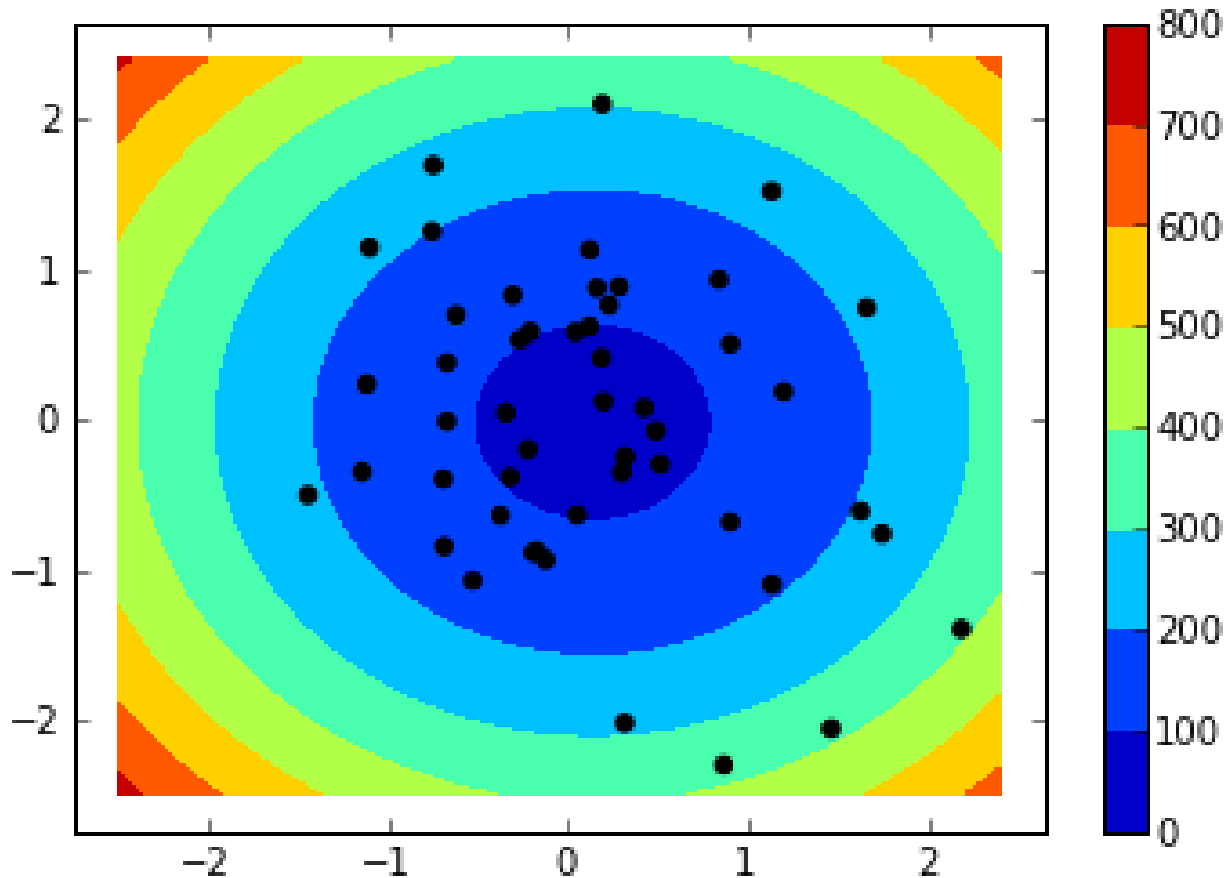
Propose an *analytical* as well as an *iterative* solution.

Example



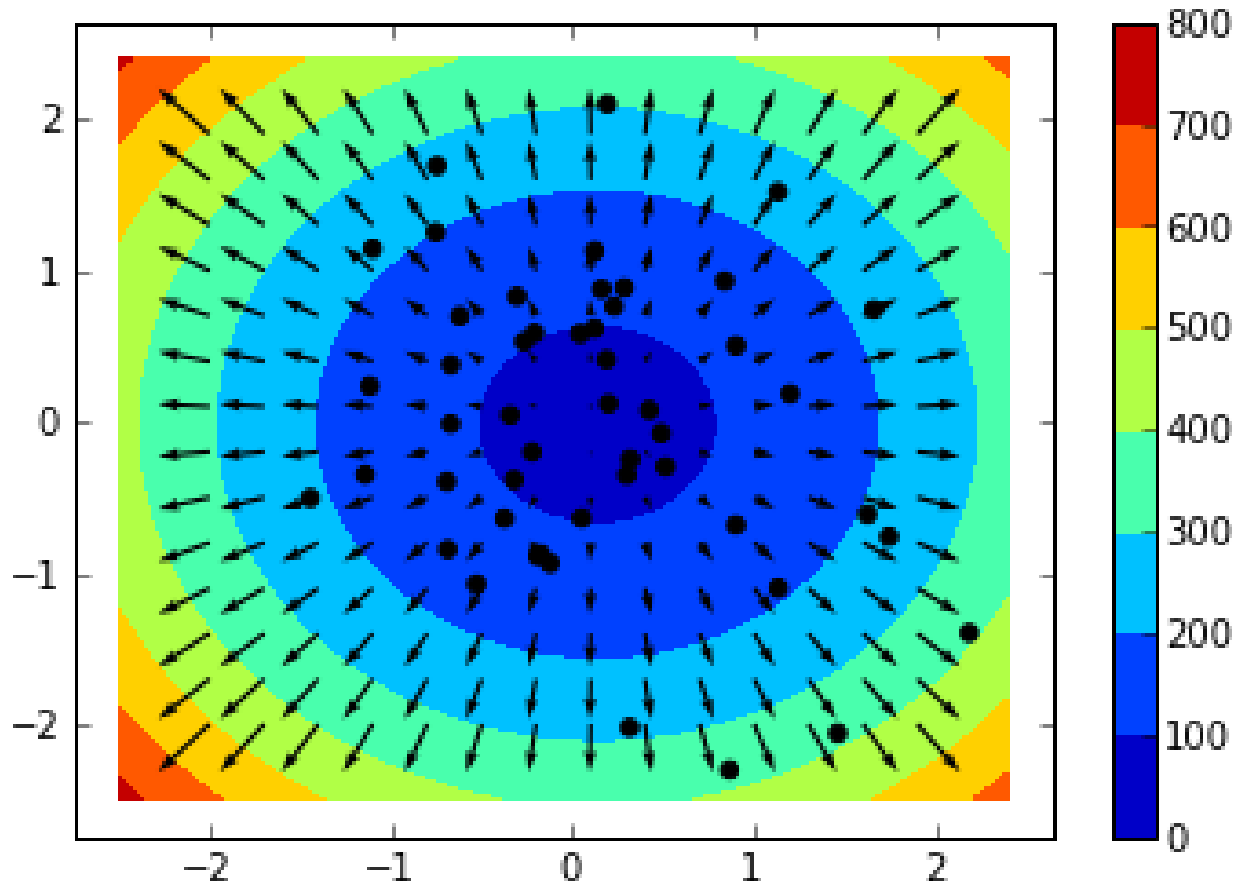
$x_1, x_2, \dots, x_{50}$

Example



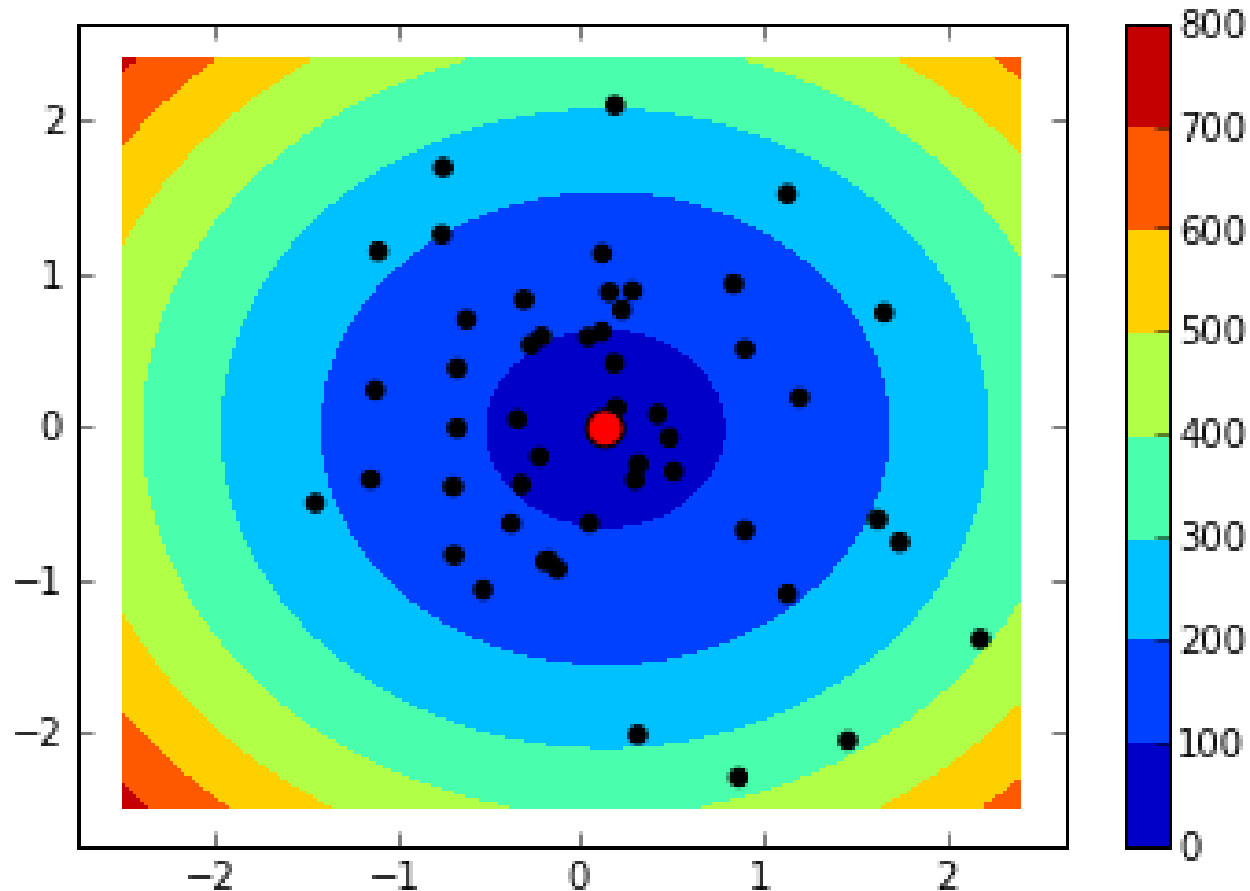
$$f(\mathbf{w}) = \sum_i \|\mathbf{x}_i - \mathbf{w}\|^2$$

Example



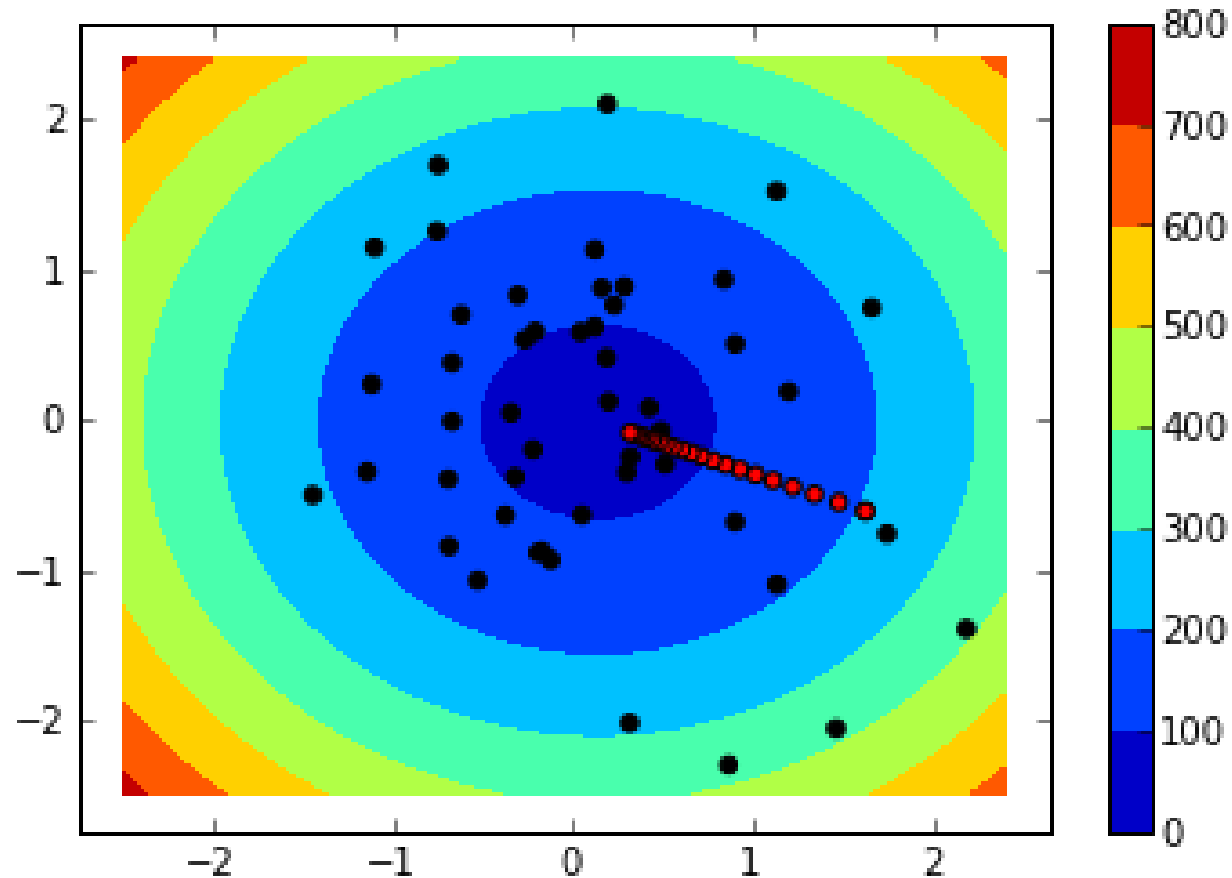
$$\nabla f(\mathbf{w}) = - \sum_i 2(\mathbf{x}_i - \mathbf{w})$$

Example



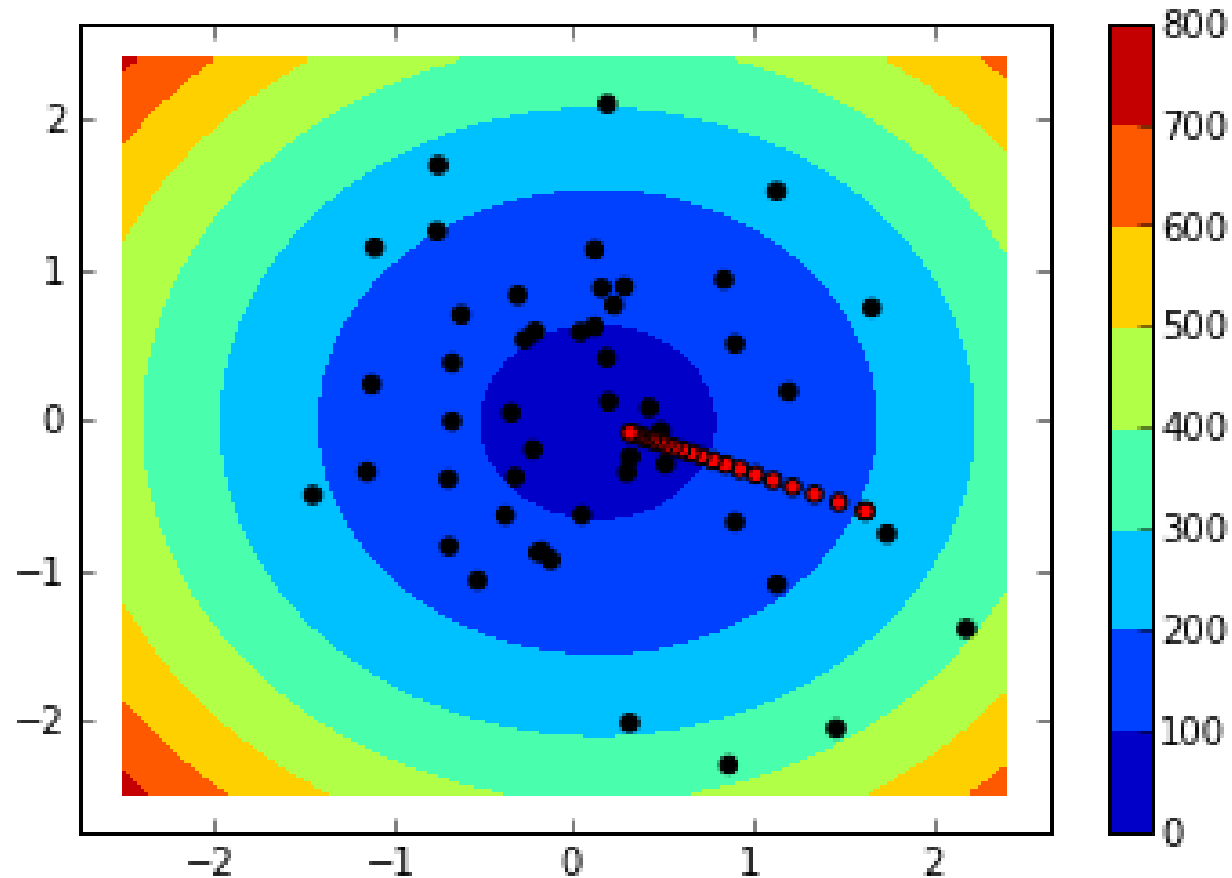
$$\nabla f(\mathbf{w}) = \mathbf{0}$$

Example



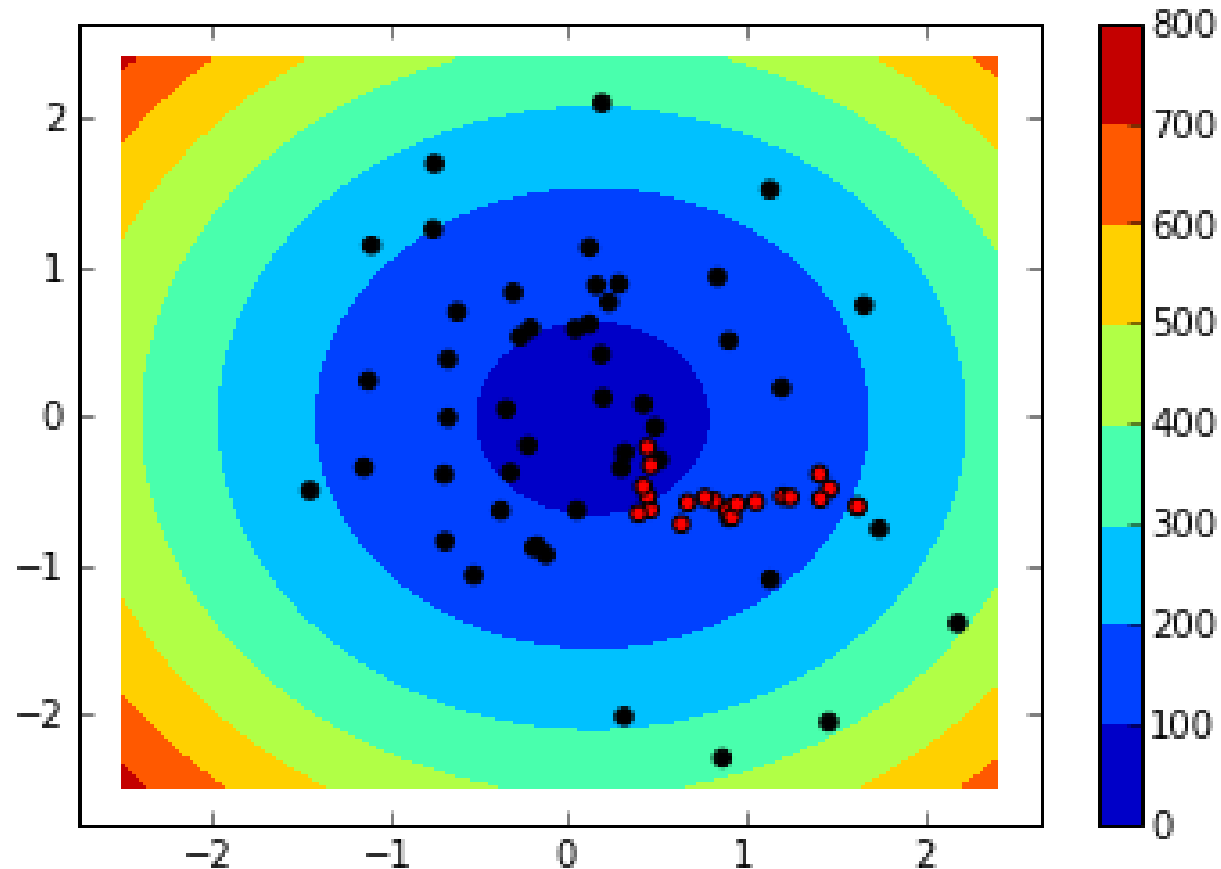
$$\Delta \mathbf{w} = -\mu \nabla f(\mathbf{w})$$

Example

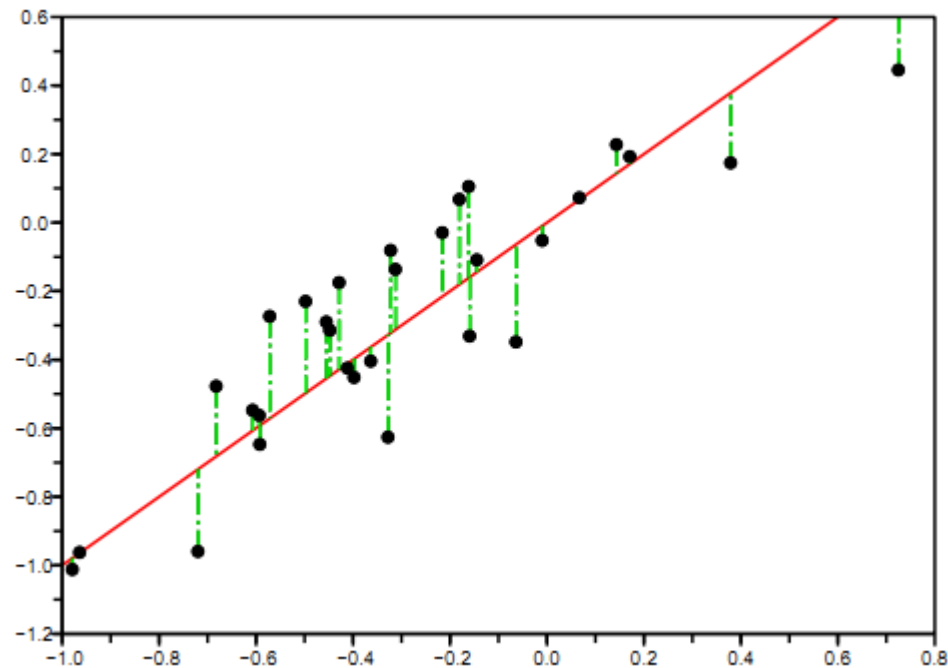


$$\Delta \mathbf{w} = \mu \sum_i 2(\mathbf{x}_i - \mathbf{w})$$

Example



$$\Delta \mathbf{w} = \mu \cdot 2(\mathbf{x}_i - \mathbf{w})$$



Linear Regression

Konstantin Tretyakov

<http://kt.era.ee>

**IFISS Summer
School 2014**

STACC

Software Technology and
Applications Competence Center



Supervised Learning

► Let X and Y be some sets.

► Let there be a training dataset:

$$D = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\}$$
$$x_i \in X, y_i \in Y$$

Supervised Learning

► Let X and Y be some sets.

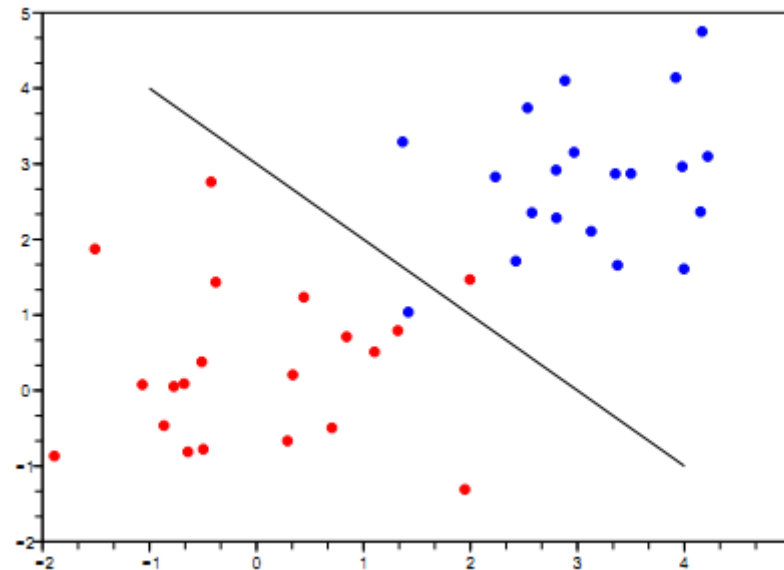
► Let there be a training dataset:

$$D = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\}$$
$$x_i \in X, y_i \in Y$$

► **Supervised learning:**

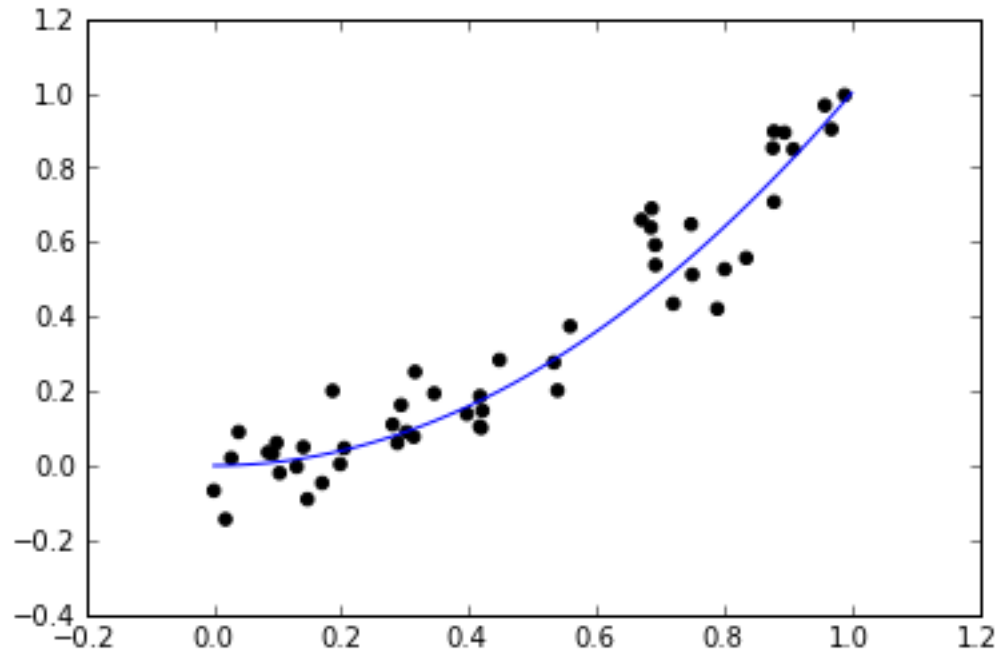
Find a function $f: X \rightarrow Y$,
generalizing the dependency
present in the data.

Classification



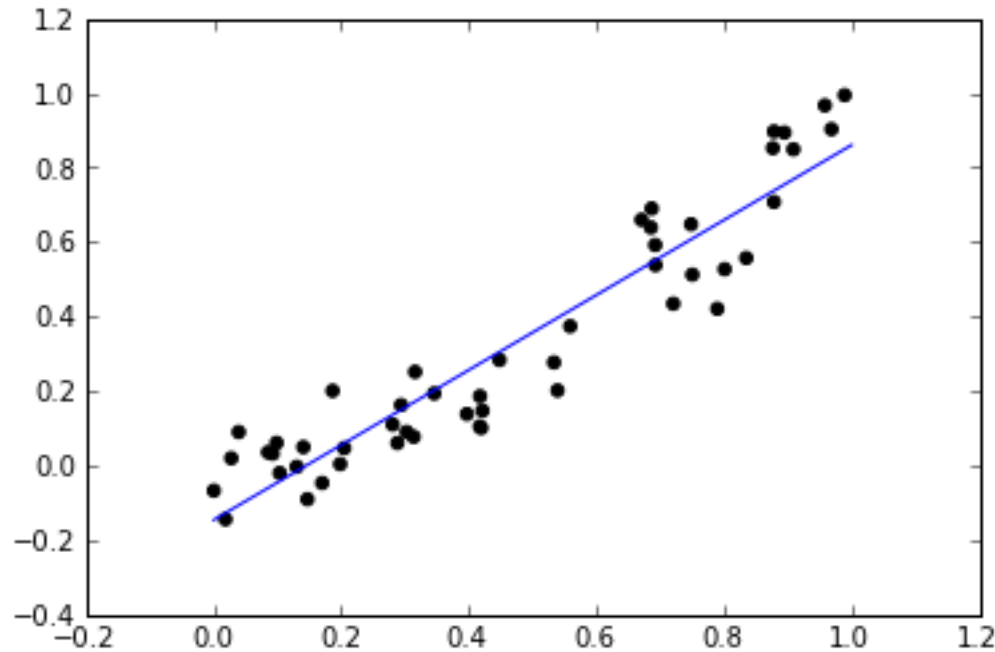
- ▶ $X = \mathbb{R}^2$, $Y = \{\text{blue}, \text{red}\}$
- ▶ $D = \{((1.3, 0.8), \text{red}), ((2.5, 2.3), \text{blue}), \dots\}$
- ▶ $f(x_1, x_2) = \text{if } (x_1 + x_2) > 3 \text{ then blue else red}$

Regression



- ▶ $X = \mathbb{R}, \quad Y = \mathbb{R}$
- ▶ $D = \{(0.5, 0.26), (0.43, 0.08), \dots\}$
- ▶ $f(x) = x^2$

Linear Regression



- ▶ $X = \mathbb{R}, \quad Y = \mathbb{R}$
- ▶ $D = \{(0.5, 0.26), (0.43, 0.08), \dots\}$
- ▶ $f(x) = -0.14 + 1.01 x$

Linear Regression

$$X = \mathbb{R}^m, \quad Y = \mathbb{R}$$

$$f(x_1, \dots, x_m) = w_0 + w_1 x_1 + \dots + w_m x_m$$

Linear Regression

$$X = \mathbb{R}^m, \quad Y = \mathbb{R}$$

$$f(x_1, \dots, x_m) = w_0 + w_1 x_1 + \dots + w_m x_m$$

$$f(x_1, \dots, x_m) = w_0 + \langle \mathbf{w}, \mathbf{x} \rangle$$

Inner product

Linear Regression

$$X = \mathbb{R}^m, \quad Y = \mathbb{R}$$

$$f(x_1, \dots, x_m) = w_0 + w_1 x_1 + \dots + w_m x_m$$

$$f(x_1, \dots, x_m) = w_0 + \langle \mathbf{w}, \mathbf{x} \rangle$$

$$f(x_1, \dots, x_m) = w_0 + (w_1, \dots, w_m) \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_m \end{pmatrix}$$

Linear Regression

$$X = \mathbb{R}^m, \quad Y = \mathbb{R}$$

$$f(x_1, \dots, x_m) = w_0 + w_1 x_1 + \dots + w_m x_m$$

$$f(x_1, \dots, x_m) = w_0 + \langle \mathbf{w}, \mathbf{x} \rangle$$

$$f(x_1, \dots, x_m) = w_0 + (w_1, \dots, w_m) \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_m \end{pmatrix}$$

$$f(x_1, \dots, x_m) = w_0 + \mathbf{w}^T \mathbf{x}$$

Linear Regression

$$f(\mathbf{x}) = w_0 + \mathbf{w}^T \mathbf{x}$$

Linear Regression

$$f(\mathbf{x}) = w_0 + \mathbf{w}^T \mathbf{x}$$

Bias term

Linear Regression

$$f(\mathbf{x}) = w_0 + \mathbf{w}^T \mathbf{x}$$

$$f(x_1, \dots, x_m) = (w_0, w_1, \dots, w_m) \begin{pmatrix} 1 \\ x_1 \\ \vdots \\ x_m \end{pmatrix}$$

Linear Regression

$$f(\mathbf{x}) = w_0 + \mathbf{w}^T \mathbf{x}$$

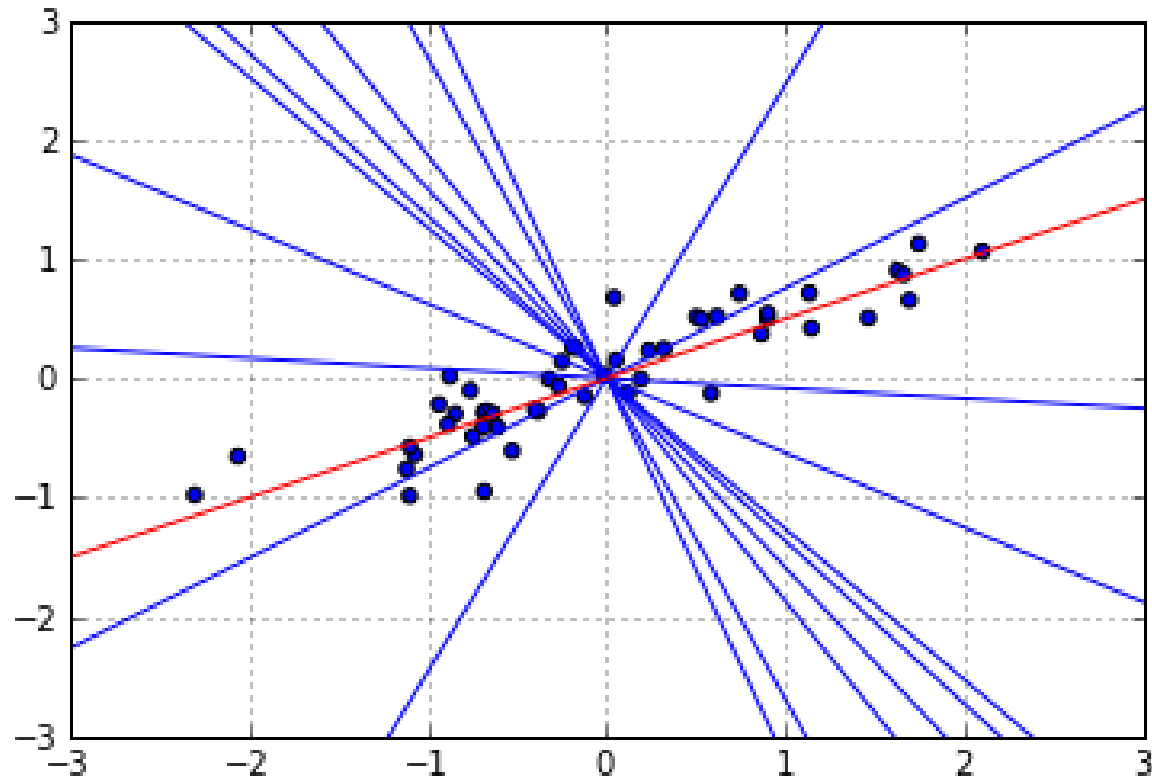
$$f(x_1, \dots, x_m) = (w_0, w_1, \dots, w_m) \begin{pmatrix} 1 \\ x_1 \\ \vdots \\ x_m \end{pmatrix}$$

$$f(\mathbf{x}) = \tilde{\mathbf{w}}^T \tilde{\mathbf{x}}$$

Linear Regression

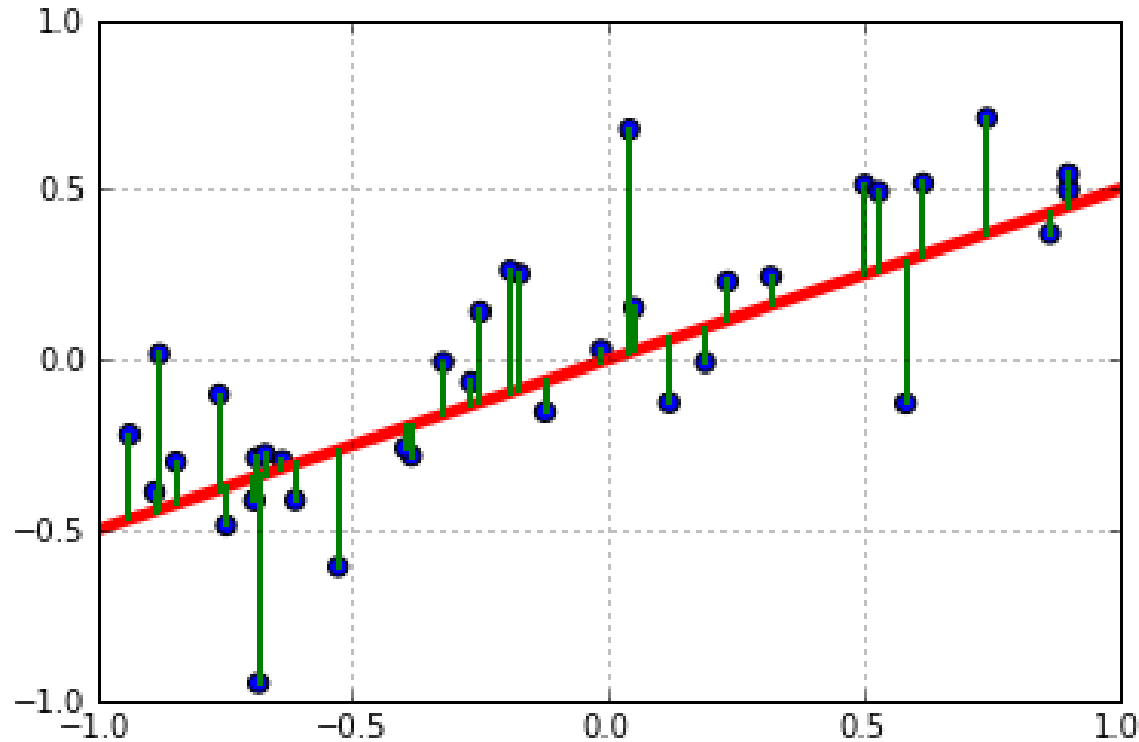
$$f(\mathbf{x}) = \mathbf{w}^T \mathbf{x}$$

Linear Regression



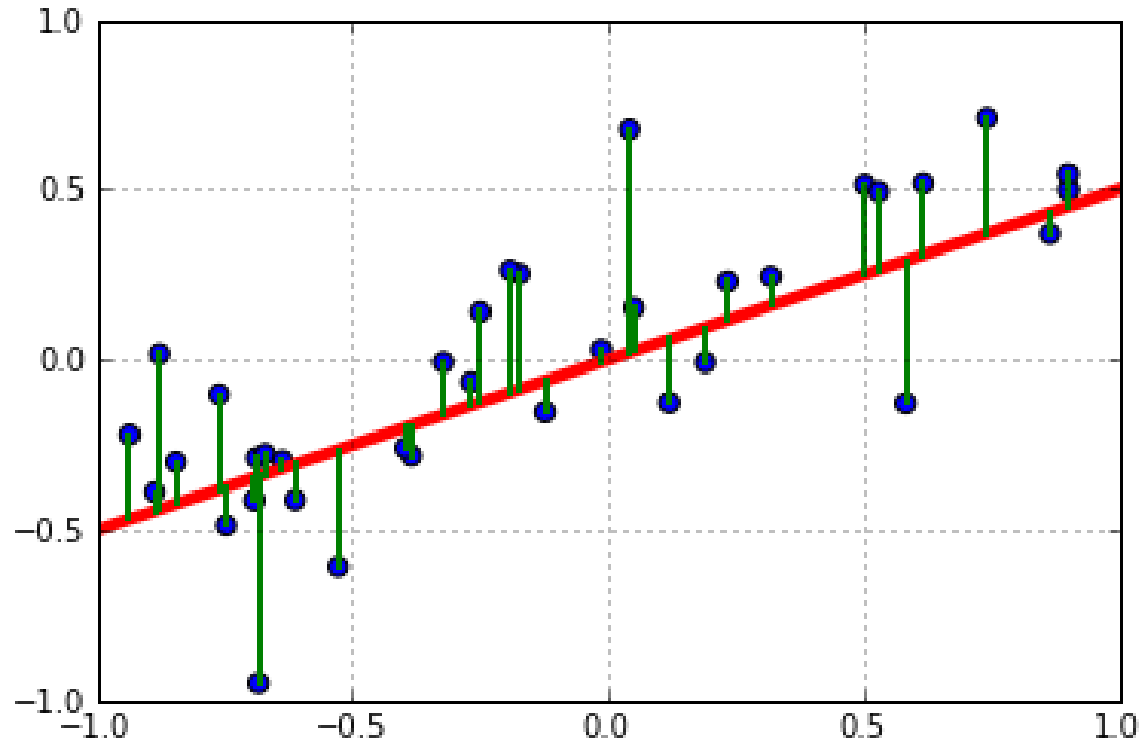
$$f(x) = wx$$

Objective Function



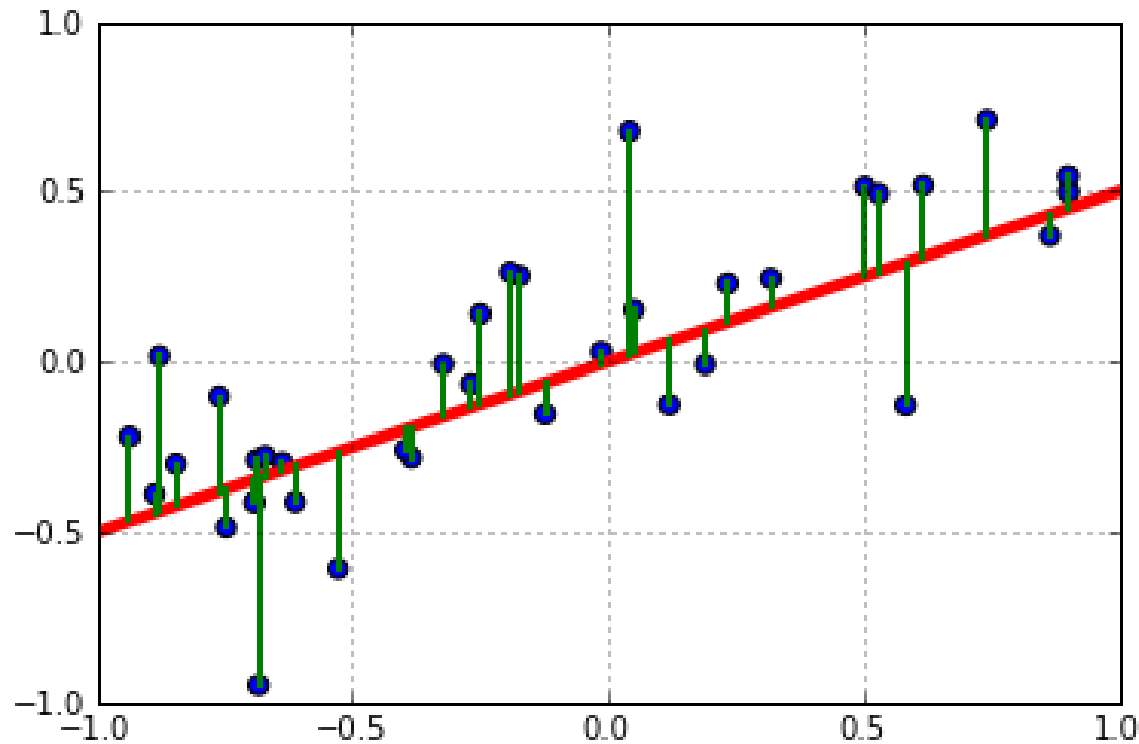
$$e_i = (f(x_i) - y_i)^2$$

Objective Function



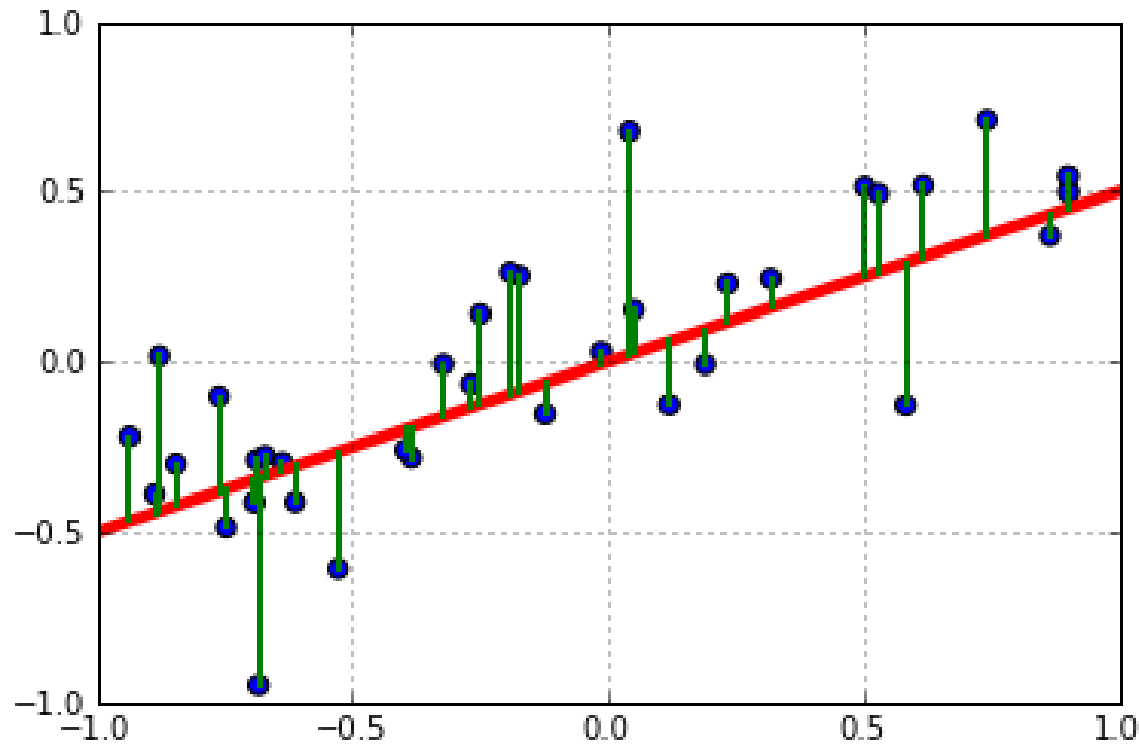
$$e_i = (\hat{y}_i - y_i)^2$$

Objective Function



$$e_i = (wx_i - y_i)^2$$

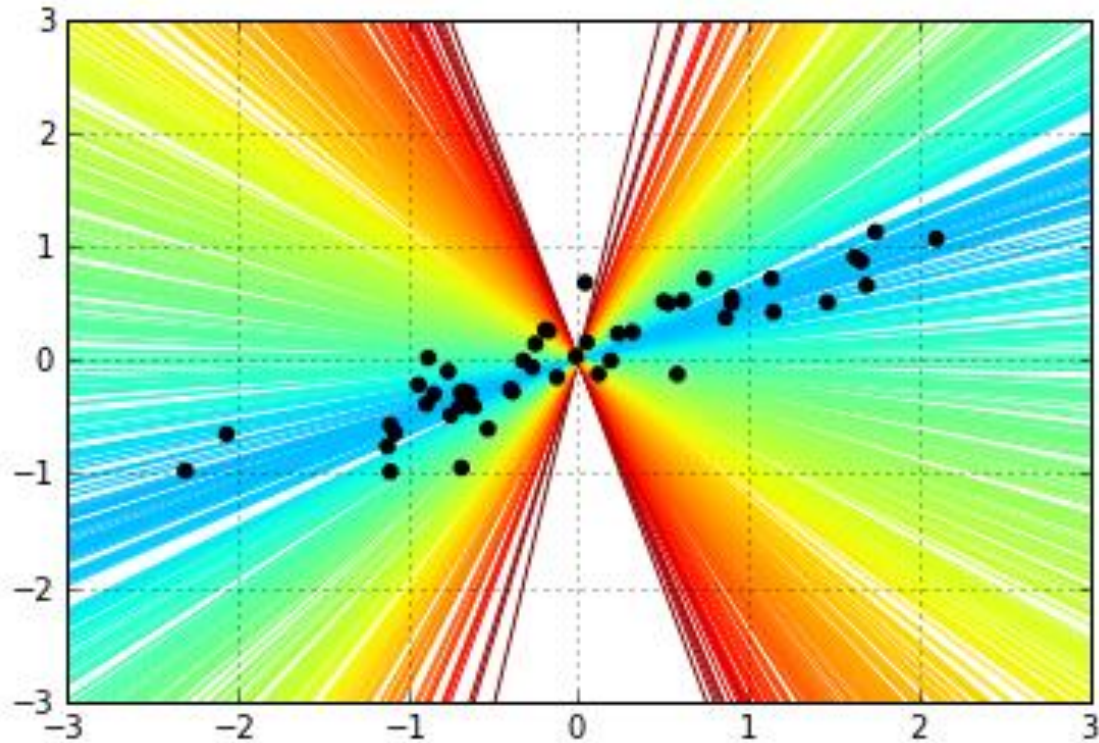
Objective Function



“OLS”

$$E_D(w) = \sum_i (wx_i - y_i)^2$$

Objective Function



$$E_D(w) = \sum_i (w x_i - y_i)^2$$

Objective Function

$$E_D(\mathbf{w}) = \sum_i (\mathbf{w}^T \mathbf{x}_i - y_i)^2$$

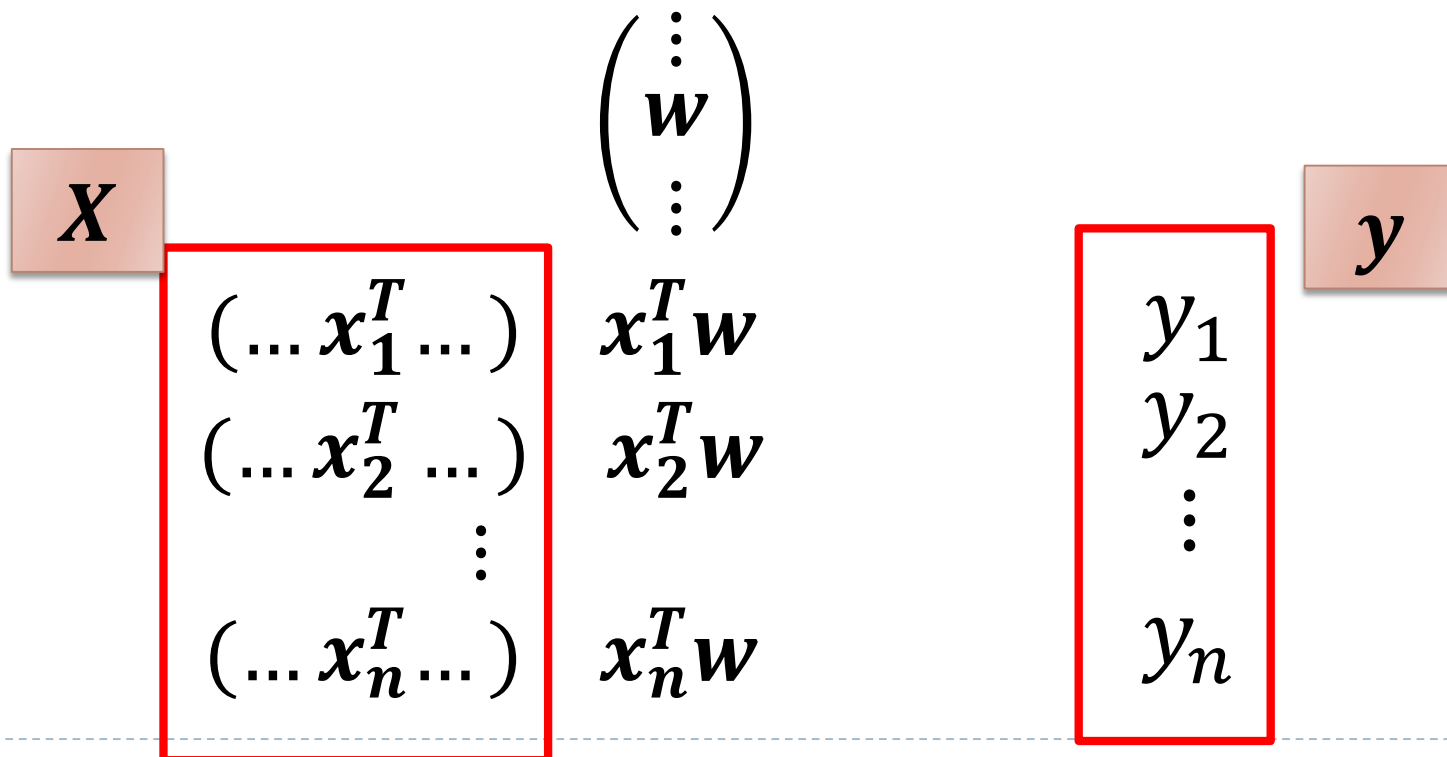
Objective Function

$$E_D(\mathbf{w}) = \sum_i (\mathbf{x}_i^T \mathbf{w} - y_i)^2$$

$$\begin{array}{ccc}
 & \begin{pmatrix} \vdots \\ \mathbf{w} \\ \vdots \end{pmatrix} & \\
 \begin{pmatrix} \dots \mathbf{x}_1^T \dots \end{pmatrix} & \mathbf{x}_1^T \mathbf{w} & y_1 \\
 \begin{pmatrix} \dots \mathbf{x}_2^T \dots \end{pmatrix} & \mathbf{x}_2^T \mathbf{w} & y_2 \\
 & \vdots & \vdots \\
 \begin{pmatrix} \dots \mathbf{x}_n^T \dots \end{pmatrix} & \mathbf{x}_n^T \mathbf{w} & y_n
 \end{array}$$

Objective Function

$$E_D(\mathbf{w}) = \sum_i (\mathbf{x}_i^T \mathbf{w} - y_i)^2$$



Objective Function

$$E_D(\mathbf{w}) = \sum_i (\mathbf{x}_i^T \mathbf{w} - y_i)^2$$

$$\mathbf{X}\mathbf{w} - \mathbf{y}$$

Objective Function

$$E_D(\mathbf{w}) = \sum_i (\mathbf{x}_i^T \mathbf{w} - y_i)^2$$

$$\|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2$$

Objective Function

$$E_D(\mathbf{w}) = \sum_i (\mathbf{x}_i^T \mathbf{w} - y_i)^2$$

$$\|X\mathbf{w} - \mathbf{y}\|^2$$

$$X\mathbf{w} \approx \mathbf{y}$$

Objective Function

$$E_D(\mathbf{w}) = \sum_i (\mathbf{x}_i^T \mathbf{w} - y_i)^2$$

$$\|X\mathbf{w} - \mathbf{y}\|^2$$

$$\mathbf{w} \approx X^{-1}\mathbf{y}?$$

$$X\mathbf{w} \approx \mathbf{y}$$

Objective Function

$$E_D(\mathbf{w}) = \sum_i (\mathbf{x}_i^T \mathbf{w} - y_i)^2$$

$$\|X\mathbf{w} - \mathbf{y}\|^2$$

$$\mathbf{w} = X^+ \mathbf{y} !$$

$$X\mathbf{w} \approx \mathbf{y}$$

Optimization

$$\operatorname{argmin}_{\mathbf{w}} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2$$

Optimization

$$\operatorname{argmin}_{\mathbf{w}} \frac{1}{2} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2$$

Optimization

$$\operatorname{argmin}_{\mathbf{w}} \frac{1}{2} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2$$

$$\|\mathbf{a}\|^2 = \mathbf{a}^T \mathbf{a}$$

Optimization

$$\operatorname{argmin}_{\mathbf{w}} \frac{1}{2} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2$$

$$\frac{1}{2} (\mathbf{X}\mathbf{w} - \mathbf{y})^T (\mathbf{X}\mathbf{w} - \mathbf{y})$$

$$(\mathbf{a} + \mathbf{b})^T = (\mathbf{a}^T + \mathbf{b}^T)$$

Optimization

$$\frac{1}{2}(\mathbf{w}^T \mathbf{X}^T - \mathbf{y}^T)(\mathbf{X}\mathbf{w} - \mathbf{y})$$

$$a(b + c) = ab + ac$$

$$\frac{1}{2}(\mathbf{w}^T \mathbf{X}^T \mathbf{X}\mathbf{w} - \mathbf{y}^T \mathbf{X}\mathbf{w} - \mathbf{w}^T \mathbf{X}^T \mathbf{y} + \mathbf{y}^T \mathbf{y})$$

Optimization

$$\frac{1}{2}(\mathbf{w}^T \mathbf{X}^T - \mathbf{y}^T)(\mathbf{X}\mathbf{w} - \mathbf{y})$$

$$\frac{1}{2}(\mathbf{w}^T \mathbf{X}^T \mathbf{X}\mathbf{w} - \mathbf{y}^T \mathbf{X}\mathbf{w} - \mathbf{w}^T \mathbf{X}^T \mathbf{y} + \mathbf{y}^T \mathbf{y})$$

$$\mathbf{y}^T \mathbf{X}\mathbf{w} = \mathbf{w}^T \mathbf{X}^T \mathbf{y} = \text{scalar}$$

$$\frac{1}{2}(\mathbf{w}^T \mathbf{X}^T \mathbf{X}\mathbf{w} - 2\mathbf{y}^T \mathbf{X}\mathbf{w} + \mathbf{y}^T \mathbf{y})$$

Optimization

$$E(\mathbf{w}) = \frac{1}{2} (\mathbf{w}^T \mathbf{X}^T \mathbf{X} \mathbf{w} - 2 \mathbf{y}^T \mathbf{X} \mathbf{w} + \mathbf{y}^T \mathbf{y})$$

Optimization

$$E(\mathbf{w}) = \frac{1}{2} (\mathbf{w}^T \mathbf{X}^T \mathbf{X} \mathbf{w} - 2 \mathbf{y}^T \mathbf{X} \mathbf{w} + \mathbf{y}^T \mathbf{y})$$

$$\nabla(f + g) = \nabla f + \nabla g$$

$$\nabla(\mathbf{w}^T \mathbf{A} \mathbf{w}) = 2 \mathbf{A} \mathbf{w}$$

$$\nabla(\mathbf{a}^T \mathbf{w}) = \mathbf{a}$$

$$\nabla E(\mathbf{w}) = \mathbf{X}^T \mathbf{X} \mathbf{w} - \mathbf{X}^T \mathbf{y}$$

Optimization

$$E(\mathbf{w}) = \frac{1}{2} (\mathbf{w}^T \mathbf{X}^T \mathbf{X} \mathbf{w} - 2 \mathbf{y}^T \mathbf{X} \mathbf{w} + \mathbf{y}^T \mathbf{y})$$

$$\nabla E(\mathbf{w}) = \mathbf{X}^T \mathbf{X} \mathbf{w} - \mathbf{X}^T \mathbf{y}$$

$$\mathbf{0} = \mathbf{X}^T \mathbf{X} \mathbf{w} - \mathbf{X}^T \mathbf{y}$$

$$\mathbf{X}^T \mathbf{X} \mathbf{w} = \mathbf{X}^T \mathbf{y}$$

Optimization

$$E(\mathbf{w}) = \frac{1}{2} (\mathbf{w}^T \mathbf{X}^T \mathbf{X} \mathbf{w} - 2 \mathbf{y}^T \mathbf{X} \mathbf{w} + \mathbf{y}^T \mathbf{y})$$

$$\nabla E(\mathbf{w}) = \mathbf{X}^T \mathbf{X} \mathbf{w} - \mathbf{X}^T \mathbf{y}$$

$$\mathbf{0} = \mathbf{X}^T \mathbf{X} \mathbf{w} - \mathbf{X}^T \mathbf{y}$$

$$\mathbf{X}^T \mathbf{X} \mathbf{w} = \mathbf{X}^T \mathbf{y}$$

$$\mathbf{w} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

Linear Regression Solution

$$\mathbf{w} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

Linear Regression Solution

$$\mathbf{w} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

`X = matrix(X)`

`y = matrix(y)`

`w = (X.T * X).I * X.T * y`

Linear Regression Solution

$$\mathbf{w} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

Moore-Penrose
pseudoinverse

`X = matrix(X)`

`y = matrix(y)`

`w = (X.T * X).I * X.T * y`

Linear Regression Solution

$$\mathbf{w} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

$$\mathbf{w} = \mathbf{X}^+ \mathbf{y}$$

Moore-Penrose
pseudoinverse

`X = matrix(X)`

`y = matrix(y)`

`w = (X.T * X).I * X.T * y`

`w = pinv(X) * y`

Linear Regression & SKLearn

```
from sklearn.linear_model import  
    LinearRegression
```

```
model = LinearRegression()  
model.fit(X, y)
```

```
w = (model.intercept_, model.coef_)  
model.predict(X_new)
```

Stochastic Gradient Regression

$$\Delta \mathbf{w} = -\mu(\mathbf{w}^T \mathbf{x}_i - y_i) \mathbf{x}_i$$

Stochastic Gradient Regression



$$\Delta \mathbf{w} = -\mu e_i \mathbf{x}_i$$



Stochastic Gradient Regression

$$\Delta \mathbf{w} = -\mu e_i \mathbf{x}_i$$

```
from sklearn.linear_model import SGDRegressor
```

```
model = SGDRegressor(alpha=0, n_iter=30)
```

```
model.fit(X, y)
```


Polynomial Regression

Say we'd like to fit a model:

$$\begin{aligned} f(x_1, x_2) \\ = w_0 + w_1 x_1 + w_2 x_2^2 + w_3 x_1 x_2 \end{aligned}$$

Polynomial Regression

Say we'd like to fit a model:

$$\begin{aligned} f(x_1, x_2) \\ = w_0 + w_1 x_1 + w_2 x_2^2 + w_3 x_1 x_2 \end{aligned}$$

Simply transform the features and proceed as normal:

$$(x_1, x_2) \rightarrow (x_1, x_2^2, x_1 x_2)$$

Single-variable Polynomial OLS

```
# n x 1 matrix
```

```
x = matrix(...)
```

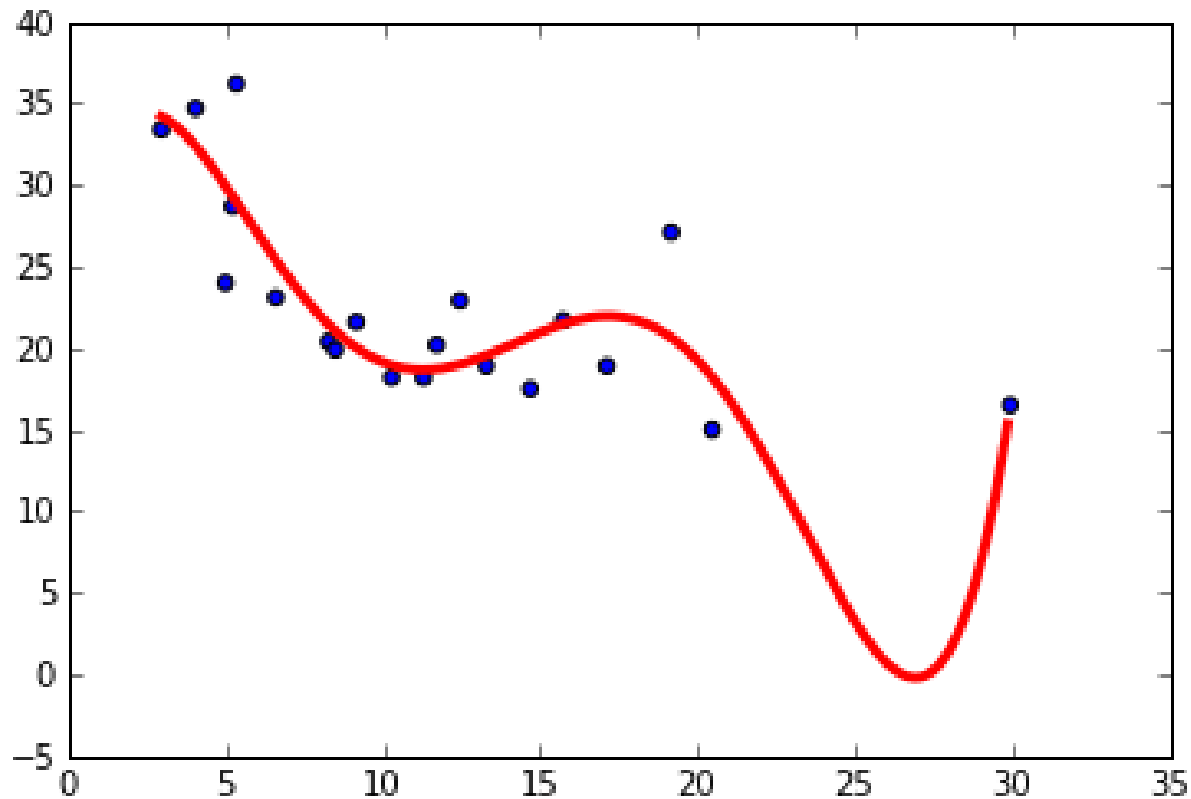
```
# Add bias & square features
```

```
X = hstack([x**0, x**1, x**2])
```

```
# Solve for w
```

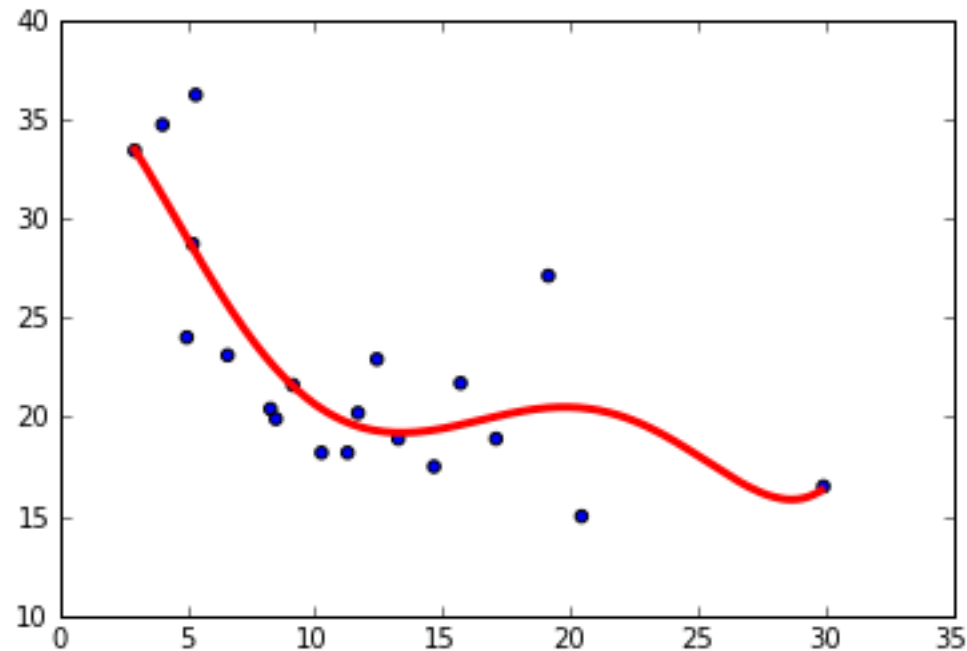
```
w = pinv(X) * y
```

Overfitting



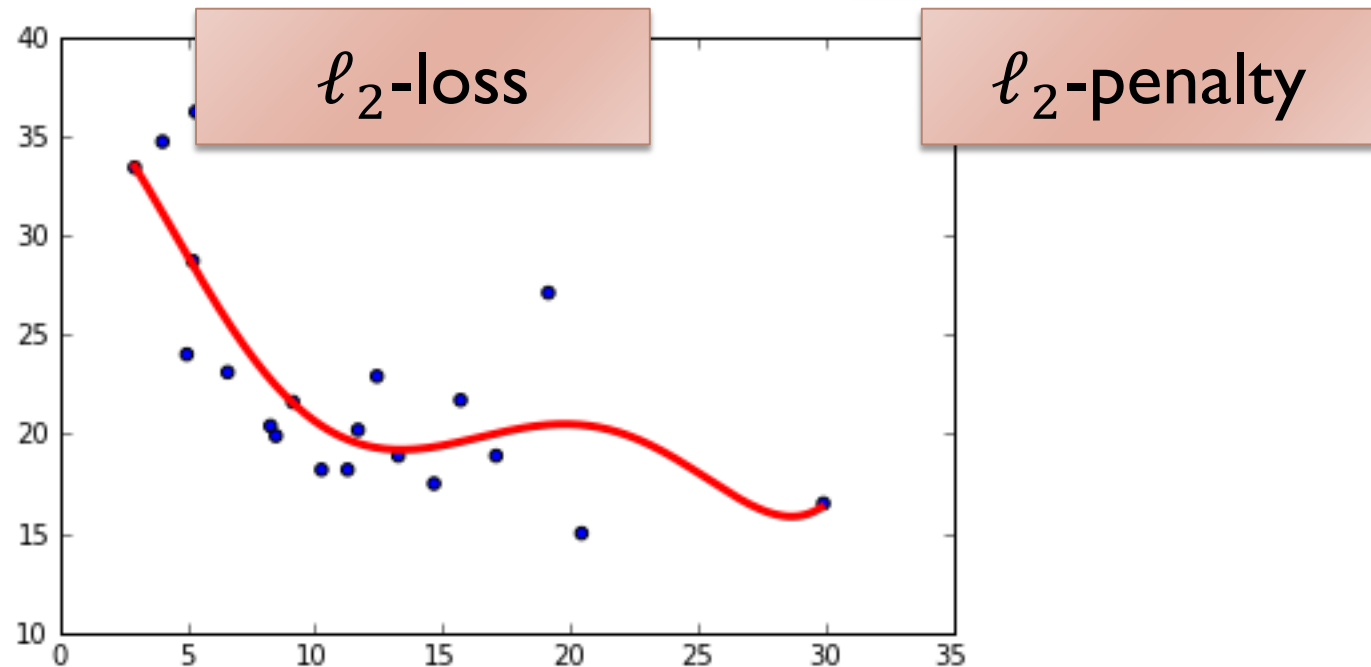
Regularization

$$E(\mathbf{w}) = \frac{1}{2} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2 + \lambda \|\mathbf{w}\|^2$$



Regularization

$$E(\mathbf{w}) = \frac{1}{2} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2 + \lambda \|\mathbf{w}\|^2$$



Regularization

$$E(\mathbf{w}) = \frac{1}{2} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2 + \lambda |\mathbf{w}|_1$$

ℓ_2 -loss

ℓ_1 -penalty

Regularization

$$E(\mathbf{w}) = \frac{1}{2} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2 + \lambda |\mathbf{w}|_0$$

ℓ_2 -loss

ℓ_0 -penalty

Regularization

$$E(\mathbf{w}) = \frac{1}{2} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2 + \lambda |\mathbf{w}|_0$$

ℓ_2 -loss

ℓ_0 -penalty

>>> SGDRegressor?

Parameters

loss : str, 'squared_loss' or 'huber' ...

...

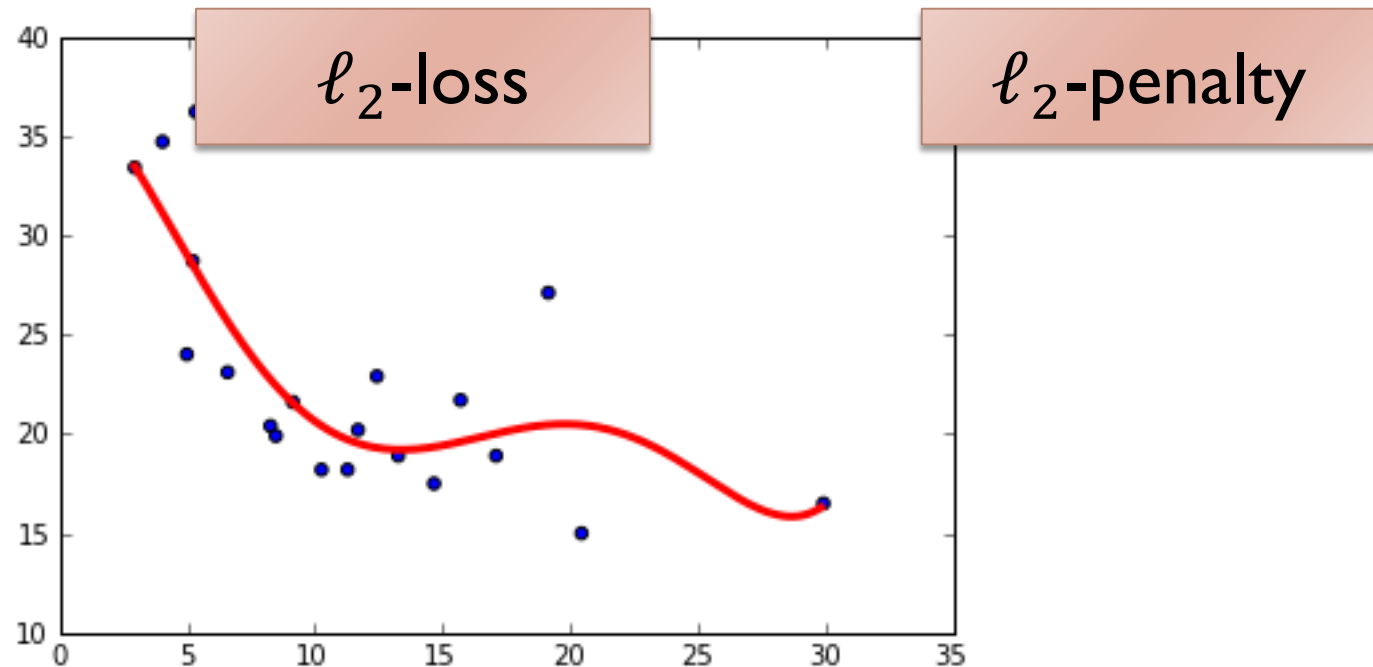
penalty : str, 'l2' or 'l1' or 'elasticnet'

...



Regularization

$$E(\mathbf{w}) = \frac{1}{2} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2 + \lambda \|\mathbf{w}\|^2$$



Ridge regression

$$\operatorname{argmin}_{\mathbf{w}} \frac{1}{2} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2 + \lambda \|\mathbf{w}\|^2$$

$$\mathbf{w} = (\mathbf{X}^T \mathbf{X} + \lambda \mathbf{I})^{-1} \mathbf{X}^T \mathbf{y}$$

$$\mathbf{I} = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix}$$

Ridge regression

$$\operatorname{argmin}_{\mathbf{w}} \frac{1}{2} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2 + \lambda \|\mathbf{w}_*\|^2$$

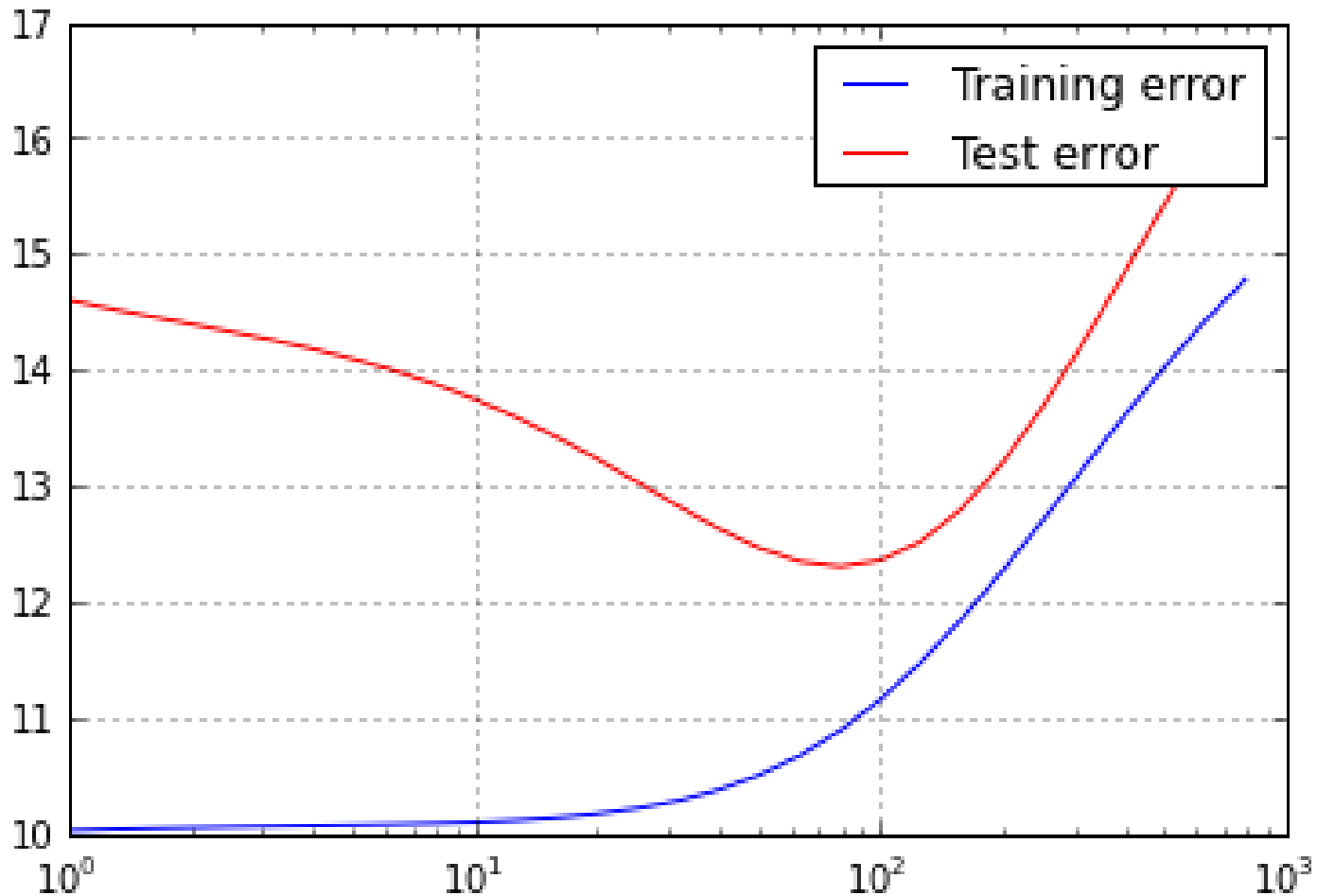
$$\mathbf{w} = (\mathbf{X}^T \mathbf{X} + \lambda \mathbf{I}_*)^{-1} \mathbf{X}^T \mathbf{y}$$

The bias term w_0 is usually **not penalized**.

$$\mathbf{I}_* = \begin{pmatrix} \boxed{0} & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 1 \end{pmatrix}$$

Derive an SGD algorithm
for Ridge Regression.

Effects of Regularization



Quiz

- ▶ OLS linear regression searches for a _____ model that has the best _____.

- ▶ Analytic solution for OLS regression:

$$\mathbf{w} = \underline{\hspace{2cm}}$$

- ▶ Stochastic gradient solution for OLS regression:

$$\Delta \mathbf{w} = \underline{\hspace{2cm}}$$

Quiz

- ▶ Large number of model parameters and/or small data may lead to _____.
- ▶ We address overfitting by _____.
- ▶ “Ridge regression” means ____-loss and ____-penalty.
- ▶ Analytic solution for Ridge regression:

$$\mathbf{w} = \underline{\hspace{2cm}}$$



Quiz

- ▶ As we increase regularization strength (i.e. increase λ), the training error _____.
- ▶ ... and the test error _____.

Questions?

