

COMS 4771 Lecture 15

1. Ordinary least squares (three perspectives)
2. Computation

REGRESSION ANALYSIS

So far, mostly looked at **classification** problems, where the label space $\mathcal{Y} = \{1, 2, \dots, K\}$ is discrete.

REGRESSION ANALYSIS

So far, mostly looked at **classification** problems, where the label space $\mathcal{Y} = \{1, 2, \dots, K\}$ is discrete.

Regression analysis considers prediction problems where label space $\mathcal{Y}$ is continuous (e.g. $[0, 1]$, $\mathbb{R}^d$, etc.)

REGRESSION ANALYSIS

So far, mostly looked at **classification** problems, where the label space $\mathcal{Y} = \{1, 2, \dots, K\}$ is discrete.

Regression analysis considers prediction problems where label space $\mathcal{Y}$ is continuous (e.g. $[0, 1]$, $\mathbb{R}^d$, etc.)

Measuring quality of a predictor $f : \mathcal{X} \rightarrow \mathcal{Y}$ using prediction error, $\Pr[f(X) \neq Y]$ doesn't make much sense in the regression context (why?)

REGRESSION ANALYSIS

So far, mostly looked at **classification** problems, where the label space $\mathcal{Y} = \{1, 2, \dots, K\}$ is discrete.

Regression analysis considers prediction problems where label space $\mathcal{Y}$ is continuous (e.g. $[0, 1]$, $\mathbb{R}^d$, etc.)

Measuring quality of a predictor $f : \mathcal{X} \rightarrow \mathcal{Y}$ using prediction error, $\Pr[f(X) \neq Y]$ doesn't make much sense in the regression context (why?)

Goal: find $f : \mathcal{X} \rightarrow \mathcal{Y}$, from some function class $\mathcal{F}$, minimizing error measured using a *loss function* $\ell : \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}$, i.e.

$$\mathbb{E}[\ell(Y, f(X))]$$

REGRESSION ANALYSIS

So far, mostly looked at **classification** problems, where the label space $\mathcal{Y} = \{1, 2, \dots, K\}$ is discrete.

Regression analysis considers prediction problems where label space $\mathcal{Y}$ is continuous (e.g. $[0, 1]$, $\mathbb{R}^d$, etc.)

Measuring quality of a predictor $f : \mathcal{X} \rightarrow \mathcal{Y}$ using prediction error, $\Pr[f(X) \neq Y]$ doesn't make much sense in the regression context (why?)

Goal: find $f : \mathcal{X} \rightarrow \mathcal{Y}$, from some function class $\mathcal{F}$, minimizing error measured using a *loss function* $\ell : \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}$, i.e.

$$\mathbb{E}[\ell(Y, f(X))]$$

E.g. if $\mathcal{Y} = \mathbb{R}$, may use $\ell(y, \hat{y}) = (y - \hat{y})^2$ ("square loss") or $\ell(y, \hat{y}) = |y - \hat{y}|$ ("absolute loss")

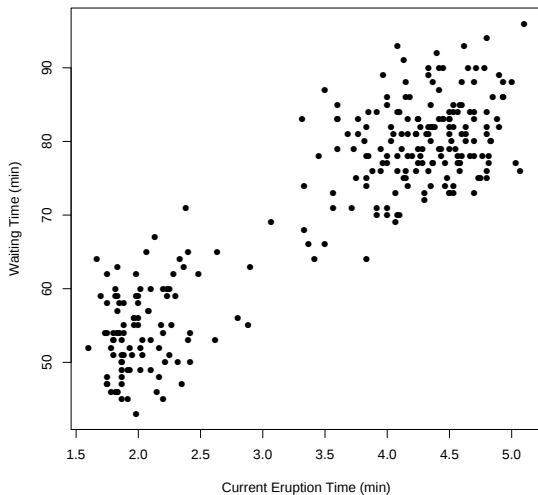
LINEAR REGRESSION (INTRODUCTION)

EXAMPLE: OLD FAITHFUL GEYSER (YELLOWSTONE)

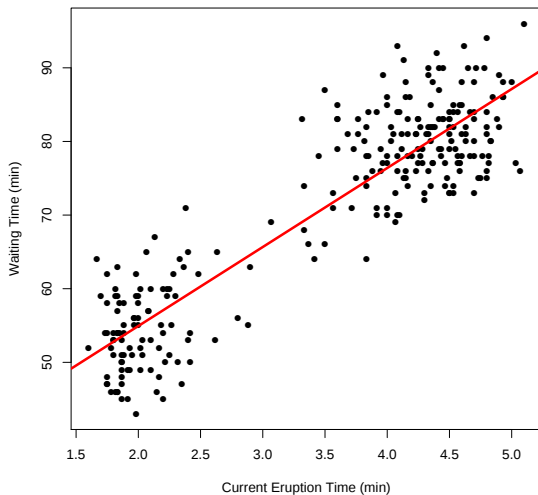


Time between eruptions seems to be related to duration of previous eruption.

EXAMPLE: OLD FAITHFUL GEYSER (YELLOWSTONE)



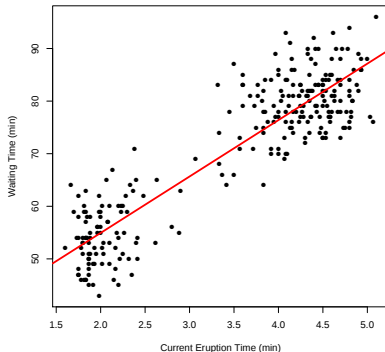
EXAMPLE: OLD FAITHFUL GEYSER (YELLOWSTONE)



EXAMPLE: OLD FAITHFUL

Linear regression model

$$(\text{wait time}) = w_0 + (\text{last duration}) \times w_1 + (\text{error})$$



MULTIVARIATE LINEAR REGRESSION

Linear regression model in $\mathbb{R}^p$

- ▶ Input variables $\mathbf{x} := (x_1, x_2, \dots, x_p)$ (“covariates”).
- ▶ Output variable y (“response”).
- ▶ Regression coefficients $\mathbf{w} := (w_1, w_2, \dots, w_p)$, intercept term w_0 .

Modeling equation:

$$y = w_0 + \langle \mathbf{x}, \mathbf{w} \rangle + \varepsilon$$

where $\varepsilon := y - (w_0 + \langle \mathbf{x}, \mathbf{w} \rangle)$ is the error term.

MULTIVARIATE LINEAR REGRESSION

Linear regression model in $\mathbb{R}^p$

- ▶ Input variables $\mathbf{x} := (x_1, x_2, \dots, x_p)$ (“covariates”).
- ▶ Output variable y (“response”).
- ▶ Regression coefficients $\mathbf{w} := (w_1, w_2, \dots, w_p)$, intercept term w_0 .

Modeling equation:

$$y = w_0 + \langle \mathbf{x}, \mathbf{w} \rangle + \varepsilon$$

where $\varepsilon := y - (w_0 + \langle \mathbf{x}, \mathbf{w} \rangle)$ is the error term.

Least squares criterion

Given pairs of input/output values, find $(w_0, \mathbf{w})$ to minimize ε^2 (on average).

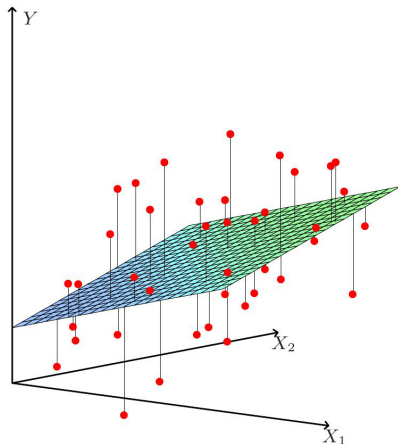
I.e. using square loss: $\varepsilon^2 = (y - \underbrace{(w_0 + \langle \mathbf{x}, \mathbf{w} \rangle)}_{\hat{y}})^2$.

LEAST SQUARES IN PICTURES

Red dots: data points.

$(w_0, w_1, w_2) \rightarrow$ affine hyperplane.

Vertical length is error.



LEAST SQUARES IN MATRIX/VECTOR FORM

Least squares criterion

Given training data

$$\mathbf{X} = \begin{bmatrix} - & \mathbf{x}^{(1)\top} & - \\ - & \mathbf{x}^{(2)\top} & - \\ & \vdots & \\ - & \mathbf{x}^{(n)\top} & - \end{bmatrix} \in \mathbb{R}^{n \times p}, \quad \mathbf{y} = \begin{bmatrix} y^{(1)} \\ y^{(2)} \\ \vdots \\ y^{(n)} \end{bmatrix},$$

find $w_0 \in \mathbb{R}$ and $\mathbf{w} \in \mathbb{R}^p$ to minimize

$$f_{\text{ls}}(\mathbf{w}) := \frac{1}{n} \sum_{i=1}^n \left(y^{(i)} - \left(w_0 + \langle \mathbf{x}^{(i)}, \mathbf{w} \rangle \right) \right)^2 = \frac{1}{n} \left\| \mathbf{y} - \begin{bmatrix} \mathbf{1} & \mathbf{X} \end{bmatrix} \begin{bmatrix} w_0 \\ \mathbf{w} \end{bmatrix} \right\|_2^2.$$

Simplification

Replace $\mathbf{X}$ with $\begin{bmatrix} \mathbf{1} & \mathbf{X} \end{bmatrix}$ and $\mathbf{w}$ with $(w_0, \mathbf{w})$, so least squares criterion is more simply written as

$$f_{\text{ls}}(\mathbf{w}) = \frac{1}{n} \left\| \mathbf{y} - \mathbf{X} \mathbf{w} \right\|_2^2.$$

LEAST SQUARES VIA CALCULUS

Least squares criterion is convex function of w ;
so suffices to find w where gradient is zero.

LEAST SQUARES VIA CALCULUS

Least squares criterion is convex function of $\mathbf{w}$;
so suffices to find $\mathbf{w}$ where gradient is zero.

Take gradient with respect to $\mathbf{w}$:

$$\nabla_{\mathbf{w}} f_{\text{ls}}(\mathbf{w}) = \nabla_{\mathbf{w}} \left\{ \frac{1}{n} \|\mathbf{y} - \mathbf{X}\mathbf{w}\|_2^2 \right\} = \frac{2}{n} \mathbf{X}^\top (\mathbf{X}\mathbf{w} - \mathbf{y}).$$

LEAST SQUARES VIA CALCULUS

Least squares criterion is convex function of $\mathbf{w}$;
so suffices to find $\mathbf{w}$ where gradient is zero.

Take gradient with respect to $\mathbf{w}$:

$$\nabla_{\mathbf{w}} f_{\text{ls}}(\mathbf{w}) = \nabla_{\mathbf{w}} \left\{ \frac{1}{n} \|\mathbf{y} - \mathbf{X}\mathbf{w}\|_2^2 \right\} = \frac{2}{n} \mathbf{X}^\top (\mathbf{X}\mathbf{w} - \mathbf{y}).$$

This is zero when

$$(\mathbf{X}^\top \mathbf{X})\mathbf{w} = \mathbf{X}^\top \mathbf{y},$$

a linear system of equations in $\mathbf{w}$ ("normal equations").

LEAST SQUARES VIA CALCULUS

Least squares criterion is convex function of $\mathbf{w}$;
so suffices to find $\mathbf{w}$ where gradient is zero.

Take gradient with respect to $\mathbf{w}$:

$$\nabla_{\mathbf{w}} f_{\text{ls}}(\mathbf{w}) = \nabla_{\mathbf{w}} \left\{ \frac{1}{n} \|\mathbf{y} - \mathbf{X}\mathbf{w}\|_2^2 \right\} = \frac{2}{n} \mathbf{X}^\top (\mathbf{X}\mathbf{w} - \mathbf{y}).$$

This is zero when

$$(\mathbf{X}^\top \mathbf{X})\mathbf{w} = \mathbf{X}^\top \mathbf{y},$$

a linear system of equations in $\mathbf{w}$ (“normal equations”).

If $\mathbf{X}^\top \mathbf{X}$ is invertible, solution is

$$\hat{\mathbf{w}}_{\text{ols}} := (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$$

(“ordinary least squares”).

COLUMN VIEW OF ORDINARY LEAST SQUARES

COLUMN VIEW OF ORDINARY LEAST SQUARES

Least squares criterion

Let $\mathbf{x}_j \in \mathbb{R}^n$ be the j -th column of $\mathbf{X} \in \mathbb{R}^{n \times p}$, so

$$\mathbf{X} = [\mathbf{x}_1 \quad \mathbf{x}_2 \quad \cdots \quad \mathbf{x}_p] .$$

COLUMN VIEW OF ORDINARY LEAST SQUARES

Least squares criterion

Let $\mathbf{x}_j \in \mathbb{R}^n$ be the j -th column of $\mathbf{X} \in \mathbb{R}^{n \times p}$, so

$$\mathbf{X} = [\mathbf{x}_1 \quad \mathbf{x}_2 \quad \cdots \quad \mathbf{x}_p] .$$

Find linear combination $\sum_{j=1}^p w_j \mathbf{x}_j$ of $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p$ so as to minimize

$$f_{\text{ls}}(\mathbf{w}) = \frac{1}{n} \|\mathbf{y} - \mathbf{X}\mathbf{w}\|_2^2 = \frac{1}{n} \left\| \mathbf{y} - \sum_{j=1}^p w_j \mathbf{x}_j \right\|_2^2 .$$

COLUMN VIEW OF ORDINARY LEAST SQUARES

Least squares criterion

Let $\mathbf{x}_j \in \mathbb{R}^n$ be the j -th column of $\mathbf{X} \in \mathbb{R}^{n \times p}$, so

$$\mathbf{X} = [\mathbf{x}_1 \quad \mathbf{x}_2 \quad \cdots \quad \mathbf{x}_p].$$

Find linear combination $\sum_{j=1}^p w_j \mathbf{x}_j$ of $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p$ so as to minimize

$$f_{\text{ls}}(\mathbf{w}) = \frac{1}{n} \|\mathbf{y} - \mathbf{X}\mathbf{w}\|_2^2 = \frac{1}{n} \left\| \mathbf{y} - \sum_{j=1}^p w_j \mathbf{x}_j \right\|_2^2.$$

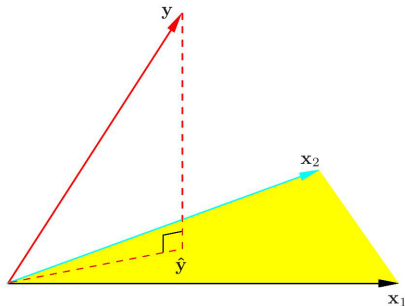
Approximation of $\mathbf{y}$ via ordinary least squares (assuming $\mathbf{X}^\top \mathbf{X}$ invertible):

$$\hat{\mathbf{y}} = \mathbf{X} \hat{\mathbf{w}}_{\text{ols}} = \sum_{j=1}^p \hat{w}_{\text{ols},j} \mathbf{x}_j.$$

COLUMN VIEW OF ORDINARY LEAST SQUARES

$\hat{\mathbf{y}} = \mathbf{X}\hat{\mathbf{w}}_{\text{ols}}$ is the orthogonal projection of $\mathbf{y}$ onto $\text{span}(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p)$:

$$\hat{\mathbf{y}} = \mathbf{X}\hat{\mathbf{w}}_{\text{ols}} = \underbrace{\mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top}_{\mathbf{\Pi}} \mathbf{y}.$$



$\mathbf{\Pi} \in \mathbb{R}^{n \times n}$ is the orthogonal projection operator for $\text{span}(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p)$.

Residual $\mathbf{r} := \mathbf{y} - \mathbf{X}\hat{\mathbf{w}}_{\text{ols}}$ is orthogonal to all $\mathbf{x}_j$.

STATISTICAL LEARNING PERSPECTIVE

Linear regression model

- ▶ Let P be a distribution over $\mathbb{R}^p \times \mathbb{R}$, and $(\mathbf{x}, y) \sim P$.

Define

$$\mathbf{w}_\star := \arg \min_{\mathbf{w} \in \mathbb{R}^p} \mathbb{E} \left[\left(y - \langle \mathbf{x}, \mathbf{w} \rangle \right)^2 \right]$$

Linear regression model

- ▶ Let P be a distribution over $\mathbb{R}^p \times \mathbb{R}$, and $(\mathbf{x}, y) \sim P$.

Define

$$\mathbf{w}_\star := \arg \min_{\mathbf{w} \in \mathbb{R}^p} \mathbb{E} \left[\left(y - \langle \mathbf{x}, \mathbf{w} \rangle \right)^2 \right]$$

- ▶ **Goal:** given i.i.d. sample S from P , find $\mathbf{w} \in \mathbb{R}^p$ so that excess mean squared error

$$\mathbb{E} \left[\left(y - \langle \mathbf{x}, \mathbf{w} \rangle \right)^2 \right] - \mathbb{E} \left[\left(y - \langle \mathbf{x}, \mathbf{w}_\star \rangle \right)^2 \right]$$

is small (and $\rightarrow 0$ as sample size $\rightarrow \infty$).

ORDINARY LEAST SQUARES

Ordinary least squares picks $\boldsymbol{w}$ to minimize empirical mean squared error based on i.i.d. sample S :

$$\hat{\boldsymbol{w}}_{\text{ols}} := \arg \min_{\boldsymbol{w} \in \mathbb{R}^p} \sum_{(\boldsymbol{x}, y) \in S} \left(y - \langle \boldsymbol{x}, \boldsymbol{w} \rangle \right)^2.$$

ORDINARY LEAST SQUARES

Ordinary least squares picks $\boldsymbol{w}$ to minimize empirical mean squared error based on i.i.d. sample S :

$$\hat{\boldsymbol{w}}_{\text{ols}} := \arg \min_{\boldsymbol{w} \in \mathbb{R}^p} \sum_{(\boldsymbol{x}, y) \in S} \left(y - \langle \boldsymbol{x}, \boldsymbol{w} \rangle \right)^2.$$

- Unconstrained convex optimization problem that happens to have “closed form” solution.

ORDINARY LEAST SQUARES

Ordinary least squares picks $\mathbf{w}$ to minimize empirical mean squared error based on i.i.d. sample S :

$$\hat{\mathbf{w}}_{\text{ols}} := \arg \min_{\mathbf{w} \in \mathbb{R}^p} \sum_{(\mathbf{x}, y) \in S} \left(y - \langle \mathbf{x}, \mathbf{w} \rangle \right)^2.$$

- ▶ Unconstrained convex optimization problem that happens to have “closed form” solution.
- ▶ Solution not uniquely defined unless $\sum_{(\mathbf{x}, y) \in S} \mathbf{x} \mathbf{x}^\top$ is invertible.

ORDINARY LEAST SQUARES

Ordinary least squares picks $\mathbf{w}$ to minimize empirical mean squared error based on i.i.d. sample S :

$$\hat{\mathbf{w}}_{\text{ols}} := \arg \min_{\mathbf{w} \in \mathbb{R}^p} \sum_{(\mathbf{x}, y) \in S} \left(y - \langle \mathbf{x}, \mathbf{w} \rangle \right)^2.$$

- ▶ Unconstrained convex optimization problem that happens to have “closed form” solution.
- ▶ Solution not uniquely defined unless $\sum_{(\mathbf{x}, y) \in S} \mathbf{x} \mathbf{x}^\top$ is invertible.
- ▶ **Predictive performance:** (with $n = |S|$)
 - ▶ $n < p$: Could be rubbish.
 - ▶ $n \geq p$: Excess mean squared error decreases at a rate of $O\left(\frac{p}{n}\right)$ (under some general conditions).

STATISTICAL ESTIMATION PERSPECTIVE

MAXIMUM LIKELIHOOD INTERPRETATION

Suppose the distribution P of $(\mathbf{x}, y)$ is such that, conditioned on $\mathbf{x}$,

$$y|\mathbf{x} \sim \mathcal{N}(\langle \mathbf{x}, \mathbf{w}_\star \rangle, \sigma^2).$$

MAXIMUM LIKELIHOOD INTERPRETATION

Suppose the distribution P of $(\mathbf{x}, y)$ is such that, conditioned on $\mathbf{x}$,

$$y|\mathbf{x} \sim \mathcal{N}(\langle \mathbf{x}, \mathbf{w}_\star \rangle, \sigma^2).$$

- **Question:** Given i.i.d. sample S , what is the MLE for $\mathbf{w}_\star$?

MAXIMUM LIKELIHOOD INTERPRETATION

Suppose the distribution P of $(\mathbf{x}, y)$ is such that, conditioned on $\mathbf{x}$,

$$y|\mathbf{x} \sim \mathcal{N}(\langle \mathbf{x}, \mathbf{w}_\star \rangle, \sigma^2).$$

- ▶ **Question:** Given i.i.d. sample S , what is the MLE for $\mathbf{w}_\star$?
- ▶ **Answer:** The ordinary least squares estimator.

Log-likelihood of $\mathbf{w}$ given S :

$$\sum_{(\mathbf{x}, y) \in S} \ln \left\{ \exp \left(-\frac{1}{2} \left(y - \langle \mathbf{x}, \mathbf{w} \rangle \right)^2 \right) \right\}$$

(plus terms that don't depend on $\mathbf{w}$).

→ maximizing likelihood $\equiv$ minimizing empirical mean squared error.

TAKING $w_\star$ SERIOUSLY

Sometimes people seriously interpret the estimated regression coefficients $\hat{w}_{\text{ols}}$.

TAKING $w_\star$ SERIOUSLY

Sometimes people seriously interpret the estimated regression coefficients $\hat{w}_{\text{ols}}$.

Example:

$\hat{w}_{\text{ols},i} = 0 \quad \longrightarrow \quad \text{variable } x_i \text{ has negligible effect on } y \text{ in } w_\star.$

$|\hat{w}_{\text{ols},i}| \gg 0 \quad \longrightarrow \quad \text{variable } x_i \text{ has significant effect on } y \text{ in } w_\star.$

TAKING $\mathbf{w}_\star$ SERIOUSLY

Sometimes people seriously interpret the estimated regression coefficients $\hat{\mathbf{w}}_{\text{ols}}$.

Example:

$$\begin{aligned}\hat{w}_{\text{ols},i} = 0 &\longrightarrow \text{variable } x_i \text{ has negligible effect on } y \text{ in } \mathbf{w}_\star. \\ |\hat{w}_{\text{ols},i}| \gg 0 &\longrightarrow \text{variable } x_i \text{ has significant effect on } y \text{ in } \mathbf{w}_\star.\end{aligned}$$

Hypothesis tests for this usually assume $y|\mathbf{x} \sim \mathcal{N}(\langle \mathbf{x}, \mathbf{w}_\star \rangle, \sigma^2)$.

COMPUTATION

Naïve computation takes $O(np^2)$ time.

FAST COMPUTATION

Naïve computation takes $O(np^2)$ time.

Hopes for speeding things up:

Naïve computation takes $O(np^2)$ time.

Hopes for speeding things up:

- ▶ Exploit sparsity or other structure in $\mathbf{X}$.

Naïve computation takes $O(np^2)$ time.

Hopes for speeding things up:

- ▶ Exploit sparsity or other structure in $\mathbf{X}$.
- ▶ Use iterative methods (e.g., *conjugate gradient* method).

- ▶ **Ordinary least squares** (three perspectives):
 1. Approximates $\mathbf{y}$ as linear combination of columns of $\mathbf{X}$.
 2. In statistical learning, excess mean squared error is $O(p/n)$.
 3. Maximum likelihood estimator under a Gaussian assumption.

- ▶ **Ordinary least squares** (three perspectives):
 1. Approximates $\mathbf{y}$ as linear combination of columns of $\mathbf{X}$.
 2. In statistical learning, excess mean squared error is $O(p/n)$.
 3. Maximum likelihood estimator under a Gaussian assumption.
- ▶ What if $p > n$?

- ▶ **Ordinary least squares** (three perspectives):
 1. Approximates $\mathbf{y}$ as linear combination of columns of $\mathbf{X}$.
 2. In statistical learning, excess mean squared error is $O(p/n)$.
 3. Maximum likelihood estimator under a Gaussian assumption.
- ▶ What if $p > n$?

Ordinary least squares no longer uniquely defined, and often ill-behaved.

- ▶ **Ordinary least squares** (three perspectives):
 1. Approximates $\mathbf{y}$ as linear combination of columns of $\mathbf{X}$.
 2. In statistical learning, excess mean squared error is $O(p/n)$.
 3. Maximum likelihood estimator under a Gaussian assumption.
- ▶ What if $p > n$?

Ordinary least squares no longer uniquely defined, and often ill-behaved.

Use regularization:

- ▶ force $\|\mathbf{w}\|_2^2$ to be small (“ridge regression”) $\rightarrow$ can kernelize this
- ▶ force $\|\mathbf{w}\|_1$ to be small (“Lasso”)
- ▶ force $\mathbf{w}$ to be sparse (“sparse regression”)

- ▶ **Ordinary least squares** (three perspectives):
 1. Approximates $\mathbf{y}$ as linear combination of columns of $\mathbf{X}$.
 2. In statistical learning, excess mean squared error is $O(p/n)$.
 3. Maximum likelihood estimator under a Gaussian assumption.
- ▶ What if $p > n$?

Ordinary least squares no longer uniquely defined, and often ill-behaved.

Use regularization:

- ▶ force $\|\mathbf{w}\|_2^2$ to be small (“ridge regression”) $\rightarrow$ can kernelize this
- ▶ force $\|\mathbf{w}\|_1$ to be small (“Lasso”)
- ▶ force $\mathbf{w}$ to be sparse (“sparse regression”)

We’ll study OLS, RR, Lasso in a simplified “fixed-design” setting.