

Jordan canonical form. let $T: V \rightarrow V$ be a linear map. Then up to change of basis T is given by a block diagonal matrix: $\begin{bmatrix} \square & & \\ & \square & \\ & & \square \end{bmatrix}$ where each block looks like $\begin{pmatrix} \lambda_i & 1 & & 0 \\ & \lambda_i & 1 & \\ & & \ddots & \ddots \\ 0 & & & \lambda_i \end{pmatrix}$ (over $\mathbb{C}$!).

Examples $\begin{bmatrix} 2 & 0 \\ 0 & 2 \end{bmatrix}$ $\begin{bmatrix} 2 & 1 \\ 0 & 2 \end{bmatrix}$ $\begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix}$ $\begin{bmatrix} 2 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 2 \end{bmatrix}$ $\begin{bmatrix} 2 & 1 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 2 \end{bmatrix}$ $\begin{bmatrix} 2 & 1 & 0 \\ 0 & 2 & 1 \\ 0 & 0 & 2 \end{bmatrix}$

A matrix is diagonalizable if its Jordan canonical form is diagonal.

"Generic" matrices have distinct eigenvalues and are diagonalizable.

Observation $A = CDC^{-1}$ $D = \begin{bmatrix} \lambda_1 & & \\ & \ddots & \\ & & \lambda_d \end{bmatrix}$

then $A^n = (CDC^{-1})^n = CDC^{-1}CDC^{-1} \dots CDC^{-1} = C D^n C^{-1} = C \begin{bmatrix} \lambda_1^n & & \\ & \ddots & \\ & & \lambda_d^n \end{bmatrix} C^{-1}$

so if λ_1 is eigenvalue with largest absolute value, then

$$\|A^n\| \sim \|A\|^n$$

Fact/Corollary $\det(A) = \text{product of eigenvalues } \lambda_1 \lambda_2 \dots \lambda_d$.

Q: how to find $\det(A)$ / eigenvalues? A : row reduction. (no scaling of rows!) $\uparrow O(d^3)$.

Useful fact If A is a symmetric matrix, then A is diagonalizable.

Proof (special case, eigenvalues all distinct).

let $v_1, \dots, v_d$ be eigenvectors for A , with corresponding eigenvalues $\lambda_1, \dots, \lambda_d$

let $C = [v_1 \ v_2 \ \dots \ v_d] \leftarrow$ matrix of eigenvectors.

claim $C^{-1}AC$ is diagonal matrix $\begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_d \end{bmatrix}$.

let $e_1 = \begin{bmatrix} 1 \\ 0 \\ \vdots \end{bmatrix}$ $e_2 = \begin{bmatrix} 0 \\ 1 \\ \vdots \end{bmatrix} \dots e_d = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{bmatrix}$.

check $Be_i = \lambda_i e_i$ for each i .

$$C^{-1} A C e_i = C^{-1} A v_i = C^{-1} \lambda_i v_i = \lambda_i C^{-1} v_i = \lambda_i e_i \quad \checkmark \quad \square$$

Dual vector spaces

Let V be a vector space, the dual vector space or dual space is

$$\begin{aligned} V^* &= \{ \text{linear maps from } V \text{ to } \mathbb{R} \} \\ &= \{ v^*: V \rightarrow \mathbb{R} \mid v^* \text{ is linear} \} \end{aligned}$$

fact V^* is a vector space:

$$\begin{aligned} 0: V^* &\rightarrow \mathbb{R} \\ v &\mapsto 0 \end{aligned}$$

operations:

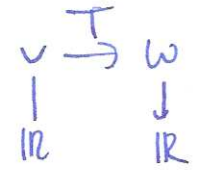
$$\left. \begin{aligned} (a^* + b^*)(v) &= a^*(v) + b^*(v) \\ (\lambda a^*)(v) &= \lambda a^*(v) \end{aligned} \right\} \begin{array}{l} \text{check} \\ \text{linear} \\ \text{maps!} \end{array}$$

$$a^*(v+w) + b^*(v+w) = a^*(v) + b^*(v) + a^*(w) + b^*(w)$$

etc.

Note
a linear map

$T: V \rightarrow W$ defines a dual map $T^*: W^* \rightarrow V^*$.



$$T^*(w^*) \in V^*$$

check T^* is linear and unique

$$T^*(w^*)(v) = w^*(T(v))$$

Basic statistical models

parametric vs non-parametric.
 $\uparrow$ $\nwarrow$
 KNN.

examples: linear regression

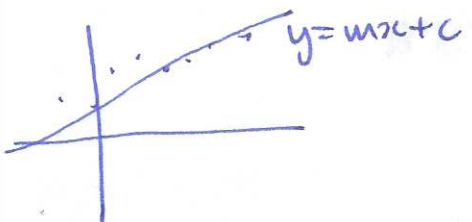
properties:

fast, get error estimates,
only works on linear data.

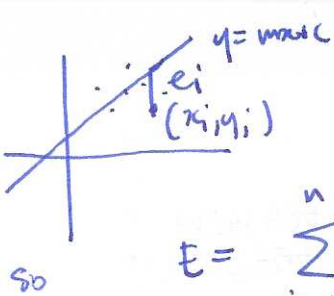
don't need to know data is linear, don't get error estimates, need to keep all data.

Review of basic linear regression

data: $(x_i, y_i) \leftarrow n \text{ points}$
 $(x_1, y_1), \dots, (x_n, y_n)$



aim: minimise least squares error LSE
 residual sum of squares RSS



$$(error)^2 = e_i^2 = (y_i - mx_i - c)^2$$

$$E = \sum_{i=1}^n (y_i - mx_i - c)^2 = \sum_{i=1}^n (y_i^2 + m^2 x_i^2 + c^2 - 2y_i m x_i - 2y_i c + 2m c x_i)$$

Q: what are best choices of m, c ? (stat people like: $\hat{\beta}_1, \hat{\beta}_0$)

1. find $\frac{\partial E}{\partial c} = \sum_{i=1}^n (2c - 2y_i + 2m x_i)$ and solve $\frac{\partial E}{\partial c} = 0$

$$\sum_{i=1}^n (2c - 2y_i + 2m x_i) = 0 \quad 2cn - 2 \sum y_i + 2m \sum x_i = 0$$

$$c = \frac{1}{n} \sum y_i - \frac{1}{n} m \sum x_i = \underbrace{\bar{y}}_{\text{average of } y\text{'s}} - m \underbrace{\bar{x}}_{\text{average of } x\text{'s}}$$

find $\frac{\partial E}{\partial m} = \sum_{i=1}^n (2m x_i^2 - 2y_i x_i + 2c x_i)$

$$= \sum_{i=1}^n (2m x_i^2 - 2y_i x_i + 2\bar{y} x_i - 2m \bar{x} x_i) \quad \text{solve } \frac{\partial E}{\partial m} = 0$$

$$m \left(\sum x_i^2 - \bar{x} \sum x_i \right) = \sum x_i y_i - \bar{y} \sum x_i$$

$$m = \frac{\left(\sum_{i=1}^n x_i y_i \right) - n \bar{x} \bar{y}}{\sum x_i^2 - n \bar{x}^2} \quad \text{fact} \quad \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

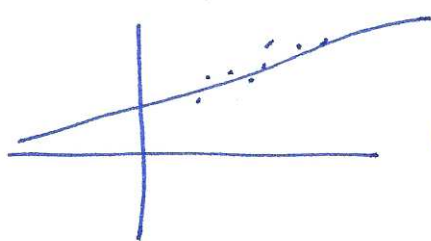
check: $\sum_{i=1}^n (x_i - \bar{x})^2 = \sum_{i=1}^n (x_i^2 - 2x_i \bar{x} + \bar{x}^2) = \sum_{i=1}^n x_i^2 - 2\bar{x}^2 n + n \bar{x}^2$

$$= \sum_{i=1}^n x_i^2 - n \bar{x}^2$$

$$\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}) = \sum (x_i y_i - \bar{x} y_i - \bar{y} x_i + \bar{x} \bar{y})$$

$$= \sum x_i y_i - n \bar{x} \bar{y} - n \bar{y} \bar{x} + n \bar{x} \bar{y} = \sum x_i y_i - n \bar{x} \bar{y} \quad \square.$$

Summary



$(x_i, y_i)_{i=1}^n$ best fit line is $y = mx + c$

with

$$m = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2} = \frac{(\sum x_i y_i) - n \bar{x} \bar{y}}{(\sum x_i^2) - n \bar{x}^2}$$

$$\text{and } c = \bar{y} - m \bar{x}.$$

note m, c depend on: $\sum x_i, \sum y_i, \sum x_i y_i, \sum x_i^2$. don't need to remember all the data.

Error estimates (assuming data is linear; errors indep of i)...

$$SE(c)^2 = \sigma^2 \left[\frac{1}{n} + \frac{\bar{x}^2}{\sum_{i=1}^n (x_i - \bar{x})^2} \right] \quad \text{Gaussian.} \quad SE(m)^2 = \frac{\sigma^2}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

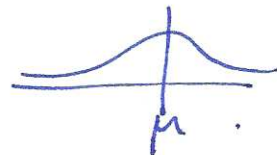
where $\sigma^2 = \text{Var}(e) \leftarrow$ variance of errors (approx by $\sum e_i^2 - \bar{e}^2$).
(in fact divide by $(n-2)$).

so can get confidence intervals:

95% confidence intervals are

$$\left. \begin{array}{l} m \pm 2 SE(m) \\ c \pm 2 SE(c) \end{array} \right\} \text{ if these values very large, probably not a good model.}$$

Recall: sample n points from a normal dist $N(\mu, \sigma^2)$.
 $x_1, \dots, x_n$



$$\bar{x} = \frac{1}{n} \sum x_i \text{ sample mean has dist } N(\mu, \frac{\sigma^2}{n})$$

Problem don't know σ in $N(\mu, \sigma^2)$. Solution: use sample variance / standard deviation approx σ by $s(x_1, \dots, x_n)$.
(should use Student's t-dist....)