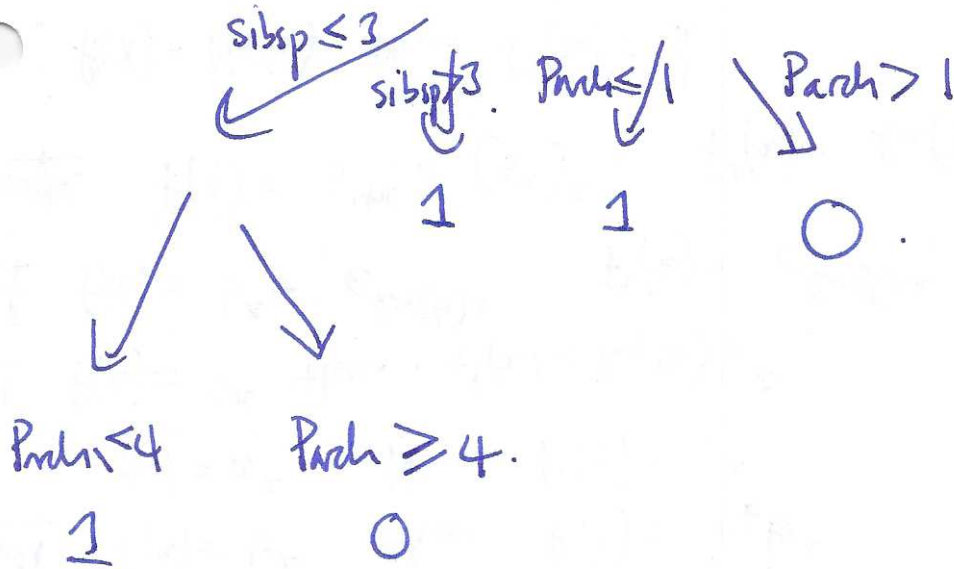


Decision trees

data
M ~~677~~ F(314)

Feb 9 ①



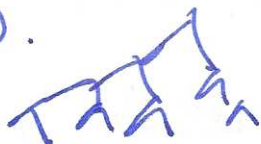
	Parch.	survived	died
0	194	153	41
1	60	46	14
2	49	30	19
3	4	3	1
4	2	0	2
5	4	1	3
6	1	0	1

choose division/output
to minimize RSS $\sum_{j=1}^J \sum_{i \in R_j} (y_i - \hat{y}_{R_j})$

Cost complexity pruning / weakest link pruning
tree $T < T_0$ $\alpha = 0$
 $T = T_0$

minimize: $\sum_{n=1}^{|T|} \sum_{i: x_i \in R_n} (y_i - \hat{y}_{R_n})^2 + \alpha |T|$

Fact as α goes from 0 to 1 get nested subtrees.



Assessing model accuracy

training data $\leadsto$ model $\leadsto$ test data.
 $X = \{x_i\}$ $f(x_i)$

do MSE
here.

mean squared error

$$MSE = \frac{1}{n} \sum_{i=1}^n (x_i - f(x_i))^2$$

avoids overfitting: given some data $X = \{x_i\}$.
 you can tune any sufficiently
 complicated model to give 100% accuracy
 on that data set in such a way it
 gives bad predictions on test data.

no test data? split training data.

jargon: variance vs bias.

sensitivity of model
to changes in data

error due to
actual data
being non-linear.

e.g. linear regression
very stable.

linear regn

low
variance

(maybe) high bias

kNN

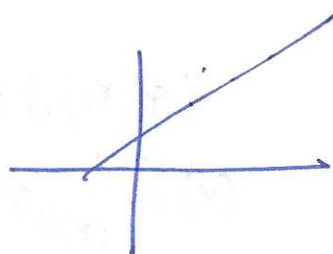
higher
variance

(maybe) low bias.

multiple linear regression

Feb 8 (3)

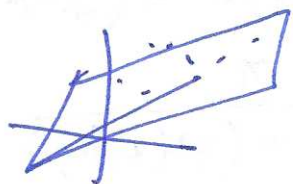
single var regnem :



(x_i, y_i)

$$y = mx + b$$

multiple var



$(x_1, x_2, \dots, x_n)$

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_n X_n$$

minimise :

$$\sum_{i=1}^n (y_i - \hat{y}_i)^2 = \sum_{i=1}^n (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_{i1} - \hat{\beta}_2 x_{i2} - \dots - \hat{\beta}_p x_{ip})^2$$
$$= \sum_{i=1}^n (y_i - [\beta][x])^2$$

for predicting z in terms of (x_i, y_i) given best fit plane.

Q: are any of the X_i useful predictors?

- do you need all of them?
- how good is the fit?
- what should we predict? How accurate is the prediction?

garden of forking paths.

- everything is probably weakly correlated
- even if things are indep, if you're at 50% significance level, 1 in 20 things are significant.