

Categorical $\equiv$ "Quantitative" (of course)

Regression Analysis used for

- Prediction (PREDICTIVE)
- Modeling (DESCRIPTIVE)
- Testing assoc. of variables to Response. (hyp. testing)

MULTIPLE LIN REG

Model Params are $\beta_0, \beta_1, \dots, \beta_n, \sigma^2$



$\Rightarrow$ UNKNOWN ALWAYS! ONLY GOD KNOWS ("GOD EQN")

$\Rightarrow$ We always provide ESTIMATES!

For (P) predictors.

Data: $\begin{bmatrix} x_1 & x_2 & x_3 & \dots & x_n \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{bmatrix}$

(rows in table)

observation of predictors

$(n \times p)$ matrix. This is easier to see for you

REST 200 M
(SORRY)!!

p2

$$Y = \beta X + \epsilon \quad (\text{there are all matrices})$$

$$\begin{bmatrix} Y_1 \\ \vdots \\ Y_n \end{bmatrix} = \begin{bmatrix} \beta_0 \\ \vdots \\ \beta_n \end{bmatrix} \cdot \begin{bmatrix} x_1 & \dots & x_{n,1} \\ \vdots & & \vdots \\ \dots & & x_{n,p} \end{bmatrix} + \begin{bmatrix} \epsilon_1 \\ \vdots \\ \epsilon_n \end{bmatrix}$$

"DESIGN MATRIX" "RESIDUAL MATRIX"

→ Order and Interactions of Regression Model.

Highest Exponent of predictors (2nd order is $\beta_0 X_1 + \beta_1 X_1^2$)
(1st order is $\beta_0 X_1 + \beta_2 X_2 + \dots$)

When you multiply two predictors ($\beta_1 X_1 X_2$)
2nd order

NOT when you multiply same pred!

Ex

1st order : $\beta_0 + \beta_1 X_1 + \beta_2 X_2 \dots$

 parallel

2nd order : $\beta_0 + \beta_1 \underline{X_1^2} + \beta_2 X_2 \dots$

 "spray"

1st order INT : $\beta_0 + \beta_1 X_1 + \beta_2 \underline{X_1 X_2} + \dots$

 curvy parallel

2nd order INT : $\beta_0 + \beta_1 \underline{X_1^2} + \beta_2 \underline{X_1 X_2}$

 curvy "spray"

Parameter Estimation

like in simple Regression, try this

NEED $\text{Min} [\sum (y_i - \hat{y}_i)^2] = \text{Min} [\sum_{i=1}^n (y_i - (\hat{\beta}_0 + \dots + \hat{\beta}_p x_{ip}))^2]$

need some linear algebra voodoo for this one

But! $\text{Min}[\sum (y_i - \hat{y}_i)^2] = (y - \underset{\substack{(\quad) \\ \text{Matrix!}}}{X\hat{\beta}})^T (y - X\hat{\beta})$

Which works out to: (via the "Orthogonal Decomposition Theorem")

$$X^T (y - \hat{y}) = X^T (y - X\hat{\beta}) = 0$$

$\Rightarrow$ We want $\hat{\beta}$! (like in simple, we want $\hat{\beta}_0, \hat{\beta}_1$)

$\Rightarrow$ Iff $X^T X$ is invertible
(IMP! Multicollinearity!)

$$\hat{\beta} = (X^T X)^{-1} X^T \cdot y$$

$$\Rightarrow \hat{y} = X\hat{\beta} = \underbrace{X(X^T X)^{-1} X^T}_{\text{"H", the Hat Matrix}} y = Hy$$

"H", the Hat Matrix

So far, only X (design matrix). What about Residuals?

Not so bad. $\hat{e} = y - \hat{y} = (I - H)y$

INTERLUDE

- You want to square shit like in simple regression?

Say $B = \begin{bmatrix} \vdots \\ \vdots \end{bmatrix} \cdot \sum B^2 = B^T \cdot B = \text{matrix} \begin{bmatrix} \vdots \\ \vdots \end{bmatrix} \begin{bmatrix} \vdots & \dots & \vdots \end{bmatrix} = \text{value}$

Sampling Dist & variance σ^2

$$\hat{\sigma}^2 \approx \text{MSE} = \frac{\hat{e}^T \hat{e}}{n(p+1)} = \frac{\hat{e}^T \hat{e}}{n-p-1} \sim \chi^2_{n-p-1} \text{ (def)}$$

???

UNDERSTAND

Why? Because we replaced p β s ... !! ???

Interpretation

"ESTIMATED"

"EXPECTED VALUE"

$\hat{\beta}_0$: EST, EV of response
when all others are zero

$\hat{\beta}_i$: EST., EV of response

① holding all others fixed!

② Change in response for
unit change in i^{th} pred. var

can be x_i

or $x_2 x_3^2, x_1 x_4$ etc!

+ Marginal & Conditional

Simple Linear Reg.

Contribution of just one
var without other

$$Y = \beta_0 + \beta_1 X + \epsilon$$

Mult L. Reg.

Assoc. of var to Resp. → INTERPRET IN
based on all other vars CONTEXT!!

→ This is about the association of vars → Resp!

→ Mult. lin. Reg. cannot be marginal!

+ Statements: Causality and Association

"This causes that"

Experimental conditions!

Change X , hold others
constant, see what
happens.

ISOLATION!

"Huh... that looks like
it's related to that..."

When you are
observing stuff

NO ISOLATION!

VARIABLES - Controlling

Explanatory

Prediction

Q1

$$Y = \beta_0 + \beta_1 X$$

population, theoretical,
unknowable!

$$\hat{Y} = \hat{\beta}_0 + \hat{\beta}_1 X$$

sample, estimate of

$$\rightarrow \boxed{\hat{\beta}_1 \text{ ESTIMATES } \beta_1}$$

(have to prove unbiasedness)

"TRUE"

→ Use Sample to make INFERENCES about population
coeffs
 $\hat{\beta}_n$ β_n

→ Regression ESTIMATES model PARAMETERS (truth)
 $\beta_0, \beta_1, \dots, \beta_n, \sigma^2$

→ Eg: $\hat{Y} = 0.89 X_1 + 0.0007 X_2$

Cannot say this is insignificant!

Eg: X_2 could be in millions! (population of people, elevation)

Why is $Y = \beta_0 + \beta_1 X_1^2 + \beta_2 \log(X_2)$ linear?

→ "Y is still regressed linearly with X_1 "

→ You put X_1^2 in Excel, R, they don't know
that it is squared!!

Q2

$$\hat{Y} = \beta_0 + \beta_1 X_1^2 + \dots + R^2$$

$$\log(\hat{Y}) = \beta_0 + \beta_1 X_1^2 + \dots + R^2$$

Cannot compare these! $\therefore$ the response is different!

$$R^2 = \frac{SSR}{SST}$$

Variance explained by model / Regression
Total variance

$$Adj(R^2) = 1 - (1 - R^2) \cdot \left(\frac{n-1}{n-k-1} \right)$$

d.f.F = n-k-1

- Adding more variables will only increase R^2
- $Adj R^2$ accounts for adding more (k) / Variables
 - ↳ Use to compare adding more k / Variables to model
 - ↳ Can be negative!

→ "Fitted Values" are $\hat{\beta}_0, \dots, \hat{\beta}_n$ in the equation just substitute 1

$$F\text{-Statistic} = \frac{MSE_{\text{Regression}}}{MSE_{\text{Error/Residuals}}} = \frac{SSR / (k-1)}{SSE / (n-k)}$$

Params
Observations
ANOVA
MULT. REG.

(Bet. Grp) (Explained Var)
(Within Grp) (Unexpl. Var)

$= \frac{SSR/p}{SSE/(n-p-1)}$

"Fitted Values" is basically $\hat{y}_i$

We assume that $\hat{e}_i \sim \hat{\epsilon}_i \sim N(0, \sigma^2)$

if this is the case, the sampling distrib of $\hat{\epsilon}_i^2$, variance, is $\chi^2_{n-2} \therefore \hat{\beta}_0$ and $\hat{\beta}_1$ replace the real deal.

(for simple regression)

Q3 / $\rightarrow$ Model scope is important! It is defined by the data used to train the model. (prediction)
 Forecasting outside of scope is EXTRAPOLATION and should be avoided! "

SIMPLE REG

Params : $\beta_0, \beta_1, \sigma^2$

Estimators : $\hat{\beta}_0, \hat{\beta}_1, \hat{\sigma}^2$

(STATISTICS

so have

Stats properties!

like samp. dist!)

is mean (Exp. val) and Var.

$\hat{\sigma}^2$: χ^2 ^{Samp. dist}
 $n-2$ dof : $\hat{\beta}_0, \hat{\beta}_1$
 replace β_0, β_1

$$\hat{\sigma}^2 = \text{MSE} = \frac{\sum \hat{e}_i^2}{n-2}$$

iff $e_i \sim N(0, \sigma^2)$

$\hat{\beta}_1$

mean, $E(\hat{\beta}_1) = \beta_1$

$\therefore$ unbiased estimator

$$\text{Var}, V(\hat{\beta}_1) = \frac{\sigma^2}{S_{xx}}$$

So, Sampling dist of $\hat{\beta}_1$ is $N(\beta_1, \sigma^2/S_{xx})$

DISCUSSION

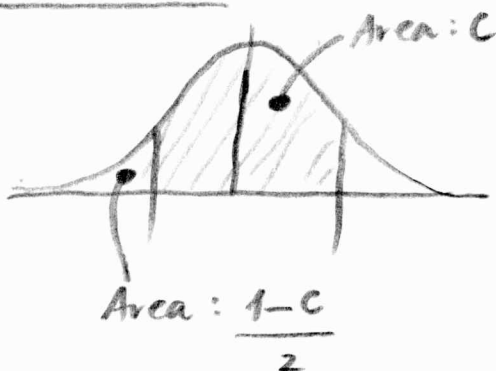
Conf. Intervals for α level :

Pop. mean = Sample mean $\pm$ Some Error

$$\text{So, } \beta_1 = \hat{\beta}_0 \pm (\text{Some Error})$$

$$= \hat{\beta}_0 \pm t_{\alpha/2, n-2} \cdot \sqrt{\frac{\text{MSE}}{S_{xx}}}$$

Error term.



$V(\hat{\beta}_1) = \sigma^2/S_{xx}$ but we don't know pop. var σ^2 , so
 replace : $\frac{\hat{\sigma}^2}{S_{xx}} = \frac{\text{MSE}}{S_{xx}}$

94/ $\Rightarrow$ "Statistically significant" means that something is different from zero! That it does something, contributes to the model.

Hypothesis testing $\xrightarrow{\text{null hyp}}$ $H_0: \beta_1 = 0$ $H_a: \beta_1 \neq 0$

like in Normal Dist, compute the "statistic" instead of "z-stat", compute "t-stat".

$$t\text{-stat} = \frac{\hat{\beta}_1 - 0}{\sqrt{\hat{\sigma}^2 / s_{xx}}}$$

} if Abs(t-stat) is large, REJECT H_0

(ie) if $t\text{-stat} > t_{\frac{\alpha}{2}, n-2}$

$$\rightarrow H_0: \beta_1 = c, \quad H_a: \beta_1 \neq c$$

same shit: $t\text{-stat} = \frac{\hat{\beta}_1 - c}{\sqrt{\hat{\sigma}^2 / s_{xx}}}$

} if Abs(t-stat) is large, REJECT H_0

(ie) if $t_{\text{stat}} > t_{\frac{\alpha}{2}, n-2}$

But! This is related to p-value.

$$p\text{-val} = 2 \times \text{Prob}(T_{n-2} > t\text{-stat})$$

$\downarrow$
 $\therefore$ sum of areas of 2 tails of dist 

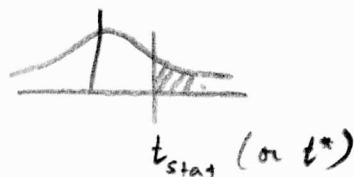
$\swarrow$ larger $t_{\text{stat}} \Rightarrow$ smaller p-val $\Rightarrow$ more evidence against null hypothesis.

($\Rightarrow$ Smaller error in denominator?)

Q6
Can test if β_1 is +ve or -ve.

$$\rightarrow \boxed{H_0: \beta_1 = 0 \quad H_A: \beta_1 < 0}$$

Only interested in one-tail of dist!



REJECT H_0 if $P(T_{n-2} \leq t_{stat}) < \alpha$

$$\rightarrow \boxed{H_0: \beta_1 = 0 \quad H_A: \beta_1 > 0}$$

one-tail! REJECT if $P(T_{n-2} > t_{stat}) > \alpha$

BACK TO INFERENCE

$$\textcircled{\hat{\beta}_0}$$

mean, $E(\hat{\beta}_0) = \beta_0$ (can prove this...)

$$\text{variance, } V(\hat{\beta}_0) = \sigma^2 \left[\frac{1}{n} + \frac{\bar{x}^2}{S_{xx}} \right]$$

↙ don't know, so replace with $\hat{\sigma}^2$ or MSE

$$\therefore \text{Var}(\hat{\beta}_0) = \text{MSE} \left[\frac{1}{n} + \frac{\bar{x}^2}{S_{xx}} \right] \quad \left(\begin{array}{l} \text{Becomes} \\ t\text{-dist w/ } n-2 \\ \text{dof!} \end{array} \right)$$

$$(\text{pop.}) \beta_0 = (\text{Sample}) \hat{\beta}_0 + (\text{Some error})$$

$$\beta_0 = \hat{\beta}_0 + t_{\frac{\alpha}{2}, n-2} \times \sqrt{\quad}$$

Can make confint with the above like w/ $\hat{\beta}_1$.

YOU NEED the normality assumption to do any of this shit!

$$\boxed{\epsilon_i \sim N(0, \sigma^2)}$$

Residual Analysis for
unrelated errors

NOT! independent errors.

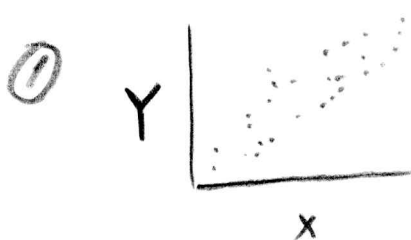
Q7/ For the line itself, Mean and Var?

$$\text{Mean} = E(\hat{Y} | x^*) = \beta_0 + \beta_1 x^* \quad \because \hat{\beta}_0, \hat{\beta}_1 \text{ are unbiased est.}$$

But Variance is diff. if in Estim mode or Pred mode
 $\therefore$ uncertainty.
 Pred mode has new var. So one more σ^2

LINEARITY
 CONST VAR
 IND.
 NORM. $\epsilon_i \sim N(0, \sigma^2)$

Assumptions & Goodness of Fit

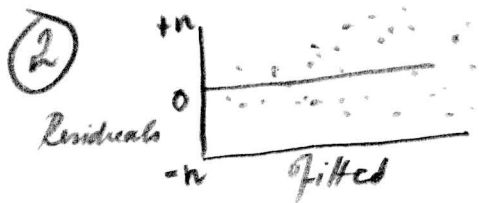


Scatterplot
 { LINEARITY
 { CONST VAR
 { INDEP

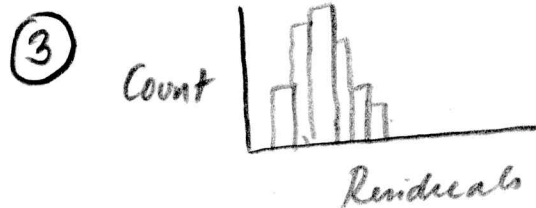
$$S_{xx} = (x_i - \bar{x})(x_i - \bar{x})$$

$$S_{xy} = (x_i - \bar{x})(y_i - \bar{y})$$

(Fit the model)

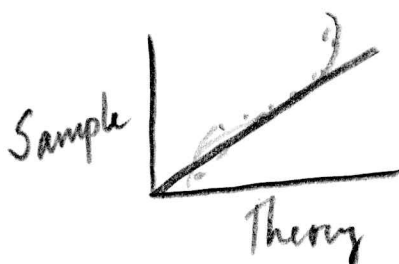


Residual Analysis (Res vs Fitted), ~~Q-Q Plot~~
 { CONST VAR
 { INDEP (uncorrelated errors!)



Histogram
 { ~~CONST VAR~~ NORMALITY

must be
 { symmetric
 { unimodal
 { no gaps



Q-Q Plot
 { NORMALITY

Q8 / ④ Transform if Problems.

NORM } Transform Resp Var, $\hat{Y}$
CONST VAR }

Box-Cox: Change $y^* = y^\lambda$

$\lambda = 1$ NO CHANGE

$\lambda = 0$ $\ln(y)$

(others...)

LINEARITY } Transform predictors x_i

ANOVA Note: - "Variance" in name but really about means (variance in means).

- About Response Data: Y_{ij} (unknown population mean)
- Model: $Y_{ij} = \mu_{ij} + \epsilon_{ij}$ → Some ~~pop~~ param to estimate.

- Assume: NORM $\epsilon_i \sim N(0, \sigma^2)$
CONS VAR $\text{Var}(\epsilon_i) = \sigma^2$
INDEP ϵ_{ij} are iid. } No linearity....
Assume Eq. Var!
Testing means.

- SST = SSE + SSR
(Total var) { Within group { Between group
N-k dof { k-1 dof
MSE = $\hat{\sigma}^2 = \frac{SSE}{N-k}$ MSR = $\frac{SSR}{k-1}$
assume equal mean
AND equal var. F-Stat = $\frac{\text{Explained}}{\text{Unexpl.}} = \frac{MSR}{MSE} = \frac{\text{Between}}{\text{Within}}$

99

You want to jack up F-stat to disprove/reject
null hypothesis: Pop. means are equal.

$$H_0: \mu_1 = \mu_2 = \dots = \mu_k$$

Simple Regression: One Quantitative

ANOVA: One or more Quantitative

Multiple Reg: One or more Quantitative & Qualitative

Marginal (vs) Conditional

Consider all other
params! when saying
something.

Causal (vs)

Experimental
Situation

$\therefore$ can change
predictors
and say something
about causation.

Association
Statement

(just observing...)

$\downarrow$ You can always
say there's an
association
(after testing w/ F, t,
hypothesis
tests
& covars)

Variable: Controlling $\rightarrow$ adjust for bias
"Default" vars

Explanatory $\rightarrow$ Explain / estimate Response

Predictive $\rightarrow$ Boost predictive power of Resp.
(even if no explanatory power)

ANOVA

Compare means ^{or treatments} $k-1$
 $H_0: \mu_1 = \mu_2 = \dots = \mu_k$ means equal! $\frac{N-k}{N-1}$

Multiple Reg

$\frac{p}{N-p-1}$ model data

$$H_0: \beta_0 = \beta_1 = \beta_2 = \dots = 0$$

the predictor coefficients are zero!

Multiple Regression: if $\lambda = 1$ no transform

$$\lambda = 0 \rightarrow \log(y) \quad (\ln(y))$$

$$\lambda = 1/2 \rightarrow \sqrt{y}$$

$$\lambda = -1 \rightarrow 1/y$$

TRUE

$$E(\epsilon_i) = 0$$

$$\text{Var}(\epsilon_i) = \sigma^2$$

$$E(\epsilon) = 0$$

$$\text{Var}(\epsilon) = \sigma^2 I$$

Identity matrix

ESTIMATE

$$E(\hat{\epsilon}_i) = 0$$

$$\text{Var}(\hat{\epsilon}_i) = \sigma^2(1 - h_{ij})$$

$$E(\hat{\epsilon}) = 0$$

$$\text{Var}(\hat{\epsilon}) = \sigma^2(I - H)$$

Shut in Hat matrix

MATRIX

^{Est.}
 $\Rightarrow$ Residuals do not have constant variance!
we have to standardize!

ie, just divide by standard dev σ (of residuals)

$$r_i = \hat{\epsilon}_i / \sigma \sqrt{1 - h_{ii}}$$

Outliers : Cannot just discard: analyze!
"Far" from the mean
LEVERAGE points $\left\{ \begin{array}{l} \text{far from mean} \\ \text{or} \\ \text{near obs. space} \end{array} \right.$
MAY OR MAY NOT affect model
INFLUENTIAL POINTS
DO affect model!

Cooks Dist: Measures how much param values
change if i^{th} obs. is removed.

General Rule if CD is $> 4/n$

or CD is > 1

take a look; analyze! do not remove!

Many outliers don't mean extreme values/
probs

Maybe a heavy-tailed dist!
correct it by transform!

- Coeff. of Determination is $R^2 = 1 - \frac{SSE}{SST} = \frac{SSR}{SST}$

"Var model explained / Total variance"

- GOF only has to do with model assumptions!
has nothing to do (does not imply) Pred. Power!

- In Mul Reg, $F\text{-stat} = \frac{MSR}{MSE} = \frac{SSR/p}{SSE/(n-p-1)}$

- Adj R^2 = Penalize for more predictors = $1 - (n-1) \left[\frac{1-R^2}{n-p-1} \right]$

- R^2 goes up w/ more p's. Doesn't mean shit.

- $(R^2) \rightarrow$ Explain Var in Response $(Adj R^2) \rightarrow$ Compare models.

- OUTLIERS: Cook's Dist. measures how much the Mul R params change when some data point is removed

Maybe problem if $CO_i > \frac{4}{n}$: investigate.

MULTICOLLINEARITY

Pearson's coeff. can help here w/ diagnosis.

linear dependence of one pred var on another.
(can happen if X_a can be estimated from X_b)

Problem? $X^T X$ must be invertible. prob. if linear dep. for OLS //

If $X^T X$ is not, SE of coeffs is ∞
SE of resp. is ∞ } theoretically.

← Practically (in R), there are "artificially large".

- R^2 doesn't change! just chill
- $\hat{\beta}_j$ is unstable (resp. changes WILDLY!)
- SE is artificially large
- F-stat is significant but t-stats are not!

VIF (Variance Inflation Factor) is used to find out.

$$VIF_j = \frac{1}{1 - R_j^2}$$

you are simply regressing X_j against the rest!

Prop ^{INCREASE} ~~Var~~ in $\text{Var}(\hat{\beta}_j)$ compared to if it were uncorrelated to something else.

Ok if $VIF < \text{Max}\left(10, \frac{1}{1 - R_{\text{model}}^2}\right)$ or 1 (lowest) (one or more... Not clear!)

High VIF $\Rightarrow$ Wider C.I.'s,
Less likely that $\hat{\beta}_j$ will be eval'd as stat sig.

VARIABLE SELECTION

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

also "Model Selection"
Art & Science!

More $X_n \rightarrow$ more BIAS less VAR

Few $X_n \rightarrow$ more VARIANCE in model.

Versus Explaining
VARIABILITY!

$\rightarrow$ We want a balance bet. BIAS and VAR

$\rightarrow$ Reduce the impact of Multicollinearity.

$\rightarrow$ PREDICTION versus EXPLANATION *

{ Want fewer
vars that can account
for variance
Easily obtainable! $\leftarrow$

{ Consider incl.
Predicting, corr.
vars $\leftarrow$

DO NOT DO THIS!

for prediction!

Other points (Esp. OVER-Interpretation)

- { Var. selection is an unresolved problem in stat "
- { No "Magic Bullet"
- $\rightarrow$ { Highly influenced by Multicollinearity! $\leftarrow$ Include one, ignore others
- { V. s. ^{can} just ignore another highly related var! $\leftarrow$ which can explain variability
- { Nothing special about the selected variables!
- { A kind of "Data Mining"
- { For Obs. Studies, CAUSALITY IS RARELY IMPLIED!!!

✓ Notation: p predictors $| X_1, X_2, \dots, X_p$
 need subset $S \subset \{1, \dots, p\}$ $|$ ^{pick} X_2, X_3, X_7
 re-index (x_j for $j \in S$) ~~$X_1, X_2, \dots, X_p$~~

"subset of vars with indices in S "

$\hat{\beta}(S) \rightarrow$ vector of est. coeffs for subset

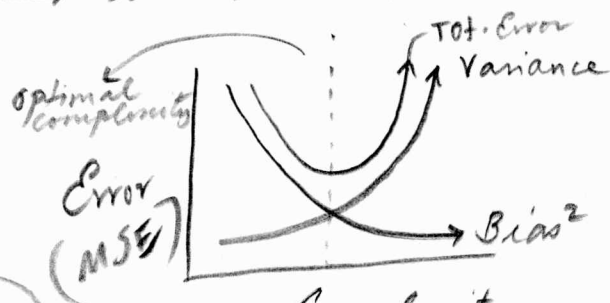
$X_S = (x_j \text{ for } j \in S)$

and $\hat{y}(S) = \hat{\beta}(S)^T X_S \hat{\beta}(S)$ } The sub-model

Some variable may be significant in full model
 but not in sub-model (sans another var)
 and vice-versa.

Also: Bias $\Rightarrow$ underfitting $\rightarrow$ use MSE to measure!
 (Bias^2)

Variance $\Rightarrow$ Overfitting



$\Rightarrow$ None of the coeffs can be signif.
BUT the overall regression can be!

Complexity
 (# of vars)
 large $\rightarrow$ small!

V3/ Prediction Risk Estimation

We want low uncertainty for new settings
 ↓
 (Future obs.)

Few $X_i \Rightarrow \uparrow \text{Bias} \downarrow \text{Var}$
 Many $X_i \Rightarrow \downarrow \text{Bias} \uparrow \text{Var}$
 very Few $X_i \Rightarrow \uparrow \text{Bias} \uparrow \text{Var}$
 ↓
 IN FUTURE OBSERV!

OPTIMIZE We want to min the Risk of making poor prediction.

How to Measure? for sub-model $R(s)$ a Risk.

$$R(s) = \sum_{i=1}^n E [\hat{Y}_i(s) - Y_i^*]^2 \quad \text{This is future obs!}$$

PRED RISK

$\hat{Y}_i(s)$ → Fitted model/response.

But don't have future obs "!", so just use values again lol

$$R_{tr}(s) = \sum_{i=1}^n [\hat{Y}_i(s) - Y_i]^2$$

TRAINING RISK

→ Uses data twice (called "snooping")

→ Always prefers the more complex model!

We obs. need to correct this...

Correction: $R_{tr}(s) + \text{Complexity Penalty}$

When model becomes too complex,
 ↓ goes down ↓ goes up

Why? We ♥ simple models!

(~~∴ prefers overfitting so that it is close to zero!~~) ~~Low Var!~~

→ "Bias upward"

1A
Higher is worse!! of course!! Complexity Penalty!

To: Conception = $R_{E'}(S) + \text{Complexity Penalty}$ | ALL THIS FOR STANDARD LIN-REG!

↳ Many of these!

① Mallow's C_p : $\frac{2|S|\hat{\sigma}^2}{n}$ → no. of predictors / model size
→ Est. variance of full model!
→ ~~Var in sub-model!~~

② AIC: $\frac{2|S|\hat{\sigma}^2}{n}$ → This is true var.
→ Replaced by est. var. in R ...
↳ Either FULL or Sub-model variance.

Bayesian

③ BIC: $\frac{|S|\hat{\sigma}^2 \log(n)}{n}$ → true var.
→ Corrects for large sample size.
→ Same as in AIC

This is used for PREDICTION more than the rest:
it penalizes complex models more!

↳ it tends to select smaller models
which may have less pred. power than AIC/Acp

④ Leave-one-out Cross-Validation

$$R_{cv}(S) \equiv \frac{1}{n} \sum_{i=1}^n E[\hat{f}_{(-i)}(S) - Y_i]$$

NOTATION! (i) means "ith value sans ith obs.!"

Can show that $R_{cv}(S) \approx R_{E'}(S) + \frac{2|S|\hat{\sigma}^2(S)}{n}$

Because of this, (use var of sub-model) AIC but of sub-model!
LOOCV penalizes complexity less than Mallow's C_p .
↓
uses FULL model!

VB For GLM, $R_{\text{err}}(s) = \frac{1}{n} \sum_{i=1}^n 2 Y_i \log \left[\frac{Y_i}{\hat{Y}_i(s)} \right]$
 { Logistic, Poisson. $+ 2 (n_i - Y_i) \log \left[\frac{n_i - Y_i}{n_i - \hat{Y}_i(s)} \right]$

(Whatever bro...)

Use AIC and BIC. $\therefore$ they are "a function of the loss-valued function"

In R, Get AIC, BIC, MCP of FULL versus SUB/Reduced Model

Lower values are obv. better!

Stepwise Regression For p params, 2^p sub-models!
 Prediction

→ "Greedy" means we take the largest or smallest jump in selection process.

if p is small, fit all
 if p is large: use heuristics/greedy algo

→ Usual caveats: ① No guarantee of "optimal" ($\therefore$ greedy)
 ② Forward model $\neq$ back model.
 (not always)

③ Forward is preferred! Smaller model $\heartsuit$

FORWARD

→ Start with no preds, Pick Criterion value C_0 (Eq. AIC)
 ④ Preds that will always be better $C_k < C_0$
 Then Add p_k (prefer p_k with smallest criterion value)

- V6 BACKWARD :- Cannot use if predictors > sample size
 $p > n$
- Computationally Expensive
 - Larger models for larger p .

REGULARIZED REGRESSION! "Penalized" Regression (mostly large models)

$$\text{Pred. Risk } R(s) = \frac{1}{n} \sum_i^n E[(\hat{Y}_i(s) - Y_i^*)]^2$$

$$= \underbrace{\text{Bias}^2[\hat{Y}_i(s)]}_{\therefore \text{of underfitting!}} + \underbrace{V(Y_i^*)}_{\text{Irreducible Error!}} + \underbrace{V(\hat{Y}_i(s))}_{\therefore \text{of overfitting!}} \rightarrow \text{MSE!}$$

This is all about Bias/Var tradeoff!

- Trying to jointly penalize β 's ("Multivariate Shrinkage")
- We kinda like prior models $\therefore$ they are smaller $\rightarrow$

In OLS, Estimate coeffs β by

$$\text{Min} \left[\sum (y_i - (\beta_0 + \beta_1 x_{i1} + \dots + \beta_p x_{ip}))^2 \right]$$

ie, sum of squared error is minimized

BUT in Reg. Reg., SSE plus penalty is minimized.

$$\text{Min} \left[\downarrow + \lambda \cdot \text{Penalty} \right]$$

$\downarrow$ Constant $\rightarrow$ determine band on β 's

✓ Important: λ balances tradeoffs!

Lack of fit
(OLS part)

Complexity
(Measure using Penalty)

it penalizes model complexity!!
large $\lambda \Rightarrow$ more penalty for complexity!

So what's the Penalty? (NO INTERCEPT!)

L_0 : Look through all sub-models, ~~pick one~~ penalize ~~for~~
~~with~~ # of non-zero β 's LOL $\rightarrow$ NORM!
not cool for large models... $\|\beta\|_0 = \#\{j: \beta_j \neq 0\}$

L_1
LASSO!

Sum abs. val of β 's $\|\beta\|_1 = \sum_{j=1}^p |\beta_j|$
Minimizing "forces" β 's to be zeros!
* Rewards sparsity!

Var Selection
 L_0, L_1

L_2
RIDGE

Sum squared β 's $\|\beta\|_2 = \sqrt{\sum_{j=1}^p \beta_j^2}$
 $\rightarrow$ Easy but does not reward sparsity!
 $\rightarrow$ Do not use for var. selection!

eg: $u = \underbrace{\{1, 0, 0, \dots, 0\}}_p$ $v = \underbrace{\{1/\sqrt{p}, 1/\sqrt{p}, \dots, 1/\sqrt{p}\}}_p$

$L_1 = 1$

$L_2 = 1$

$L_1 = \sqrt{p}$

$L_2 = 1$

Throws out
related variables!

use for
MULTICOLLINEAR!

vs Elastic-Net: Combine both!

$$Q(\beta_1, \dots, \beta_p) = \sum (OLS) + \lambda_1 \sum |\beta| + \lambda_2 \sum \beta^2$$

different!

in R, use α to turn on/off.

$$Q = OLS + \lambda [\alpha L_1 + (1-\alpha)L_2]$$

when $\alpha = 1$, only LASSO

$\alpha = 0$, only RIDGE

$0 < \alpha < 1$, ELASTIC NET

$$RSE = \sqrt{\frac{SSE}{n-p}}$$

Introduces

This is "Bias". It doesn't fit the data properly. $\Rightarrow$ Underfitting!

"Shrinkage" vs "Selection"
RIDGE LASSO

$\therefore$ squares!

Cannot go to zero
can go close...

Can go to zero
Predictions

- Ridge is good for $p > n$
- " C_p " $\rightarrow$ Complexity Penalty.

In R if P-value is 1,
model fits perfectly: "Complete Separation"

Other finer points:

Must scale predictors! (maybe responses for Lasso, EN, Ridge)

Interpret in terms of original

LASSO: No closed form, numeric algorithm

RIDGE: closed form

There are less efficient than OLS

$\therefore$ they increase bias! (Statquest video)

They inc. Bias but lower Variance.

NO
MODEL
SELECTION

LASSO will pick up to 'n' vars

usually not a problem $\because n > p$

$$\frac{n}{p} \quad \text{||} \\ \text{(preds)}$$

ElasticNet: RIDGE will dominate LASSO for multi collin...

MUST run OLS after Var. Selection!

Evaluate this stuff

Normal Regression: MSE

Logistic Regression: Classification Error

Poisson Regression: Sum of Sq. Deviances.

Mallows crpf when no control vars!

V10
 λ is determined by K-fold C.V.

for λ in $(0 \dots 1)$

① for fold in $(1 \dots k)$

fit model

then... compute/evaluate (e.g. MSE)

② Then overall error for all folds.

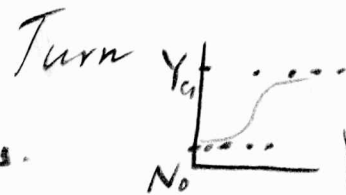
Now pick whichever λ (from ②) has smallest MSE

LOGISTIC REGRESSION

For Yes/No questions.

Problem is cannot use Normality Ass.

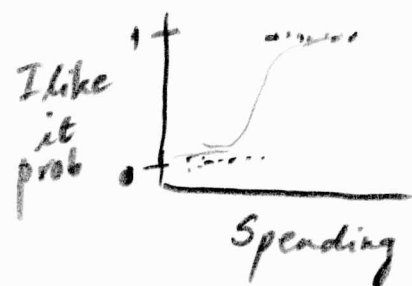
$\therefore$ Response Y is Binary...



We find that S-curve many ways to do this!

The Model Y is a binary response.

$$\left\{ \begin{array}{l} (x_{11}, x_{12}, \dots, x_{1p}), Y_1 \\ \vdots \\ (x_{n1}, x_{n2}, \dots, x_{np}), Y_n \end{array} \right\} \text{ Binary Response.}$$



so, we model Probability of success : $P(Y=1 | x_1, x_2, \dots, x_p)$

say $p = P(Y=1 | x_1, x_2, \dots, x_p)$

so get a linear model thru a non-linear link function

$$g(p) = g(p(x_1, \dots, x_p)) = \beta_0 + \beta_1 x_1 + \dots + \beta_p x_p$$

$\rightarrow$ There is no error term!

ASSUMPTIONS : ① Linearity

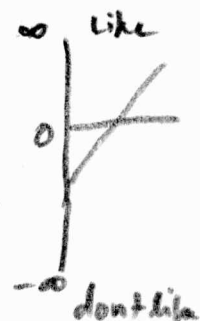
This is diff. than ordinary Regression.
Saying that whatever g is, it is a linear combo of predicting vars.

$$g(p(x_1, \dots, x_p)) = \beta_0 + \dots + \beta_p x_p$$

But g itself, does not have to be linear!

✓ ② Independence: $Y_1, \dots, Y_n$ are ind. rand. vars.

③ Logit function: Assume g is $\ln\left(\frac{p}{1-p}\right)$ | See Statquest video!
again, 'p' is prob of success!



So, $p = P(Y=1 | X_1, \dots, X_p)$

so, $\ln\left(\frac{p}{1-p}\right) = \beta_0 + \beta_1 X_1 + \dots + \beta_p X_p$

④ $p = \frac{e^{(\beta_0 + \beta_1 X_1 + \dots + \beta_p X_p)}}{1 + e^{(\beta_0 + \beta_1 X_1 + \dots + \beta_p X_p)}}$

$\frac{p}{1-p}$ is odds ... $\ln\left(\frac{p}{1-p}\right)$ is log odds (duh)

NOT the same as odds ratio!

So for a 1-unit increase, odds ratio = $\frac{e^{(\beta_0 + \beta_1 (b+1))}}{e^{\beta_0 + \beta_1 b}}$

ie, $\ln\left(\frac{p}{1-p}\right) = \beta_0 + \beta_1 X_1 + \beta_2 X_2$

$= e^{\beta_1}$

→ Interpret as the (increase or decrease) in the log of odds ratio

for a unit increase in X_1

→ THIS IS ODDS OF SUCCESS!

Nothing to do with resp. variable!

✓3/ Model Parameters Estimate are $\beta_0 \dots \beta_p$
no σ^2 because no error term!

Estimate using MLE (find the best S-curve)

- This is hard. No closed-form $\therefore \log(\text{likelihood})$ is non-linear
- Use Numerical Algo "
- Get $\hat{\beta}_0 \hat{\beta}_1 \dots \hat{\beta}_p$
- Assume Bernoulli(???)

Interlude

Bernoulli: "Will X happen or not?"

Binomial: "prob. of success after n Bernoulli trials"

↓

So in LogReg, if repetitions, use Binomial!

When doing e^{β} , $e^{-\beta} = \frac{1}{e^{\beta}}$

if val is 0.732, 27% decrease / lower than...
if val is 1.274, 27% increase / higher than...

Odds ratio is the multiplicative effect

Eg: Odds Ratio : $\frac{\text{Odds}(\text{Survive} | \text{Smoker})}{\text{Odds}(\text{Survive} | \text{Non Smoker})}$

OR = 1.46 $\Rightarrow$ 46% higher } That's it.
OR = 0.73 $\Rightarrow$ 27% lower }

Inference: MLE is pretty good for large sample sizes!

If not, Type I errors...

Not reliable otherwise !!

$$\text{sampling dist of } \hat{\beta} \approx N(\beta, V)$$

cannot get this! No error term! Get from software...
∴ no closed form expression...

So! ∴ N.D., Conf Int is Z-interval

$$(1-\alpha) \text{ Int} = \hat{\beta}_j \pm Z_{\alpha/2} \sqrt{V(\hat{\beta}_j)}$$

Hypothesis Testing: Use "Wald Test" (z-test, really)

Same shit as linear Reg BTW.

Thing is... not a t-stat ∴ of large sample size!

Small samples $\Rightarrow P(\text{Type I Error}) > \alpha$ (Sig. level)

Subset Testing: $\text{logit}(p) = \underbrace{\beta_0 + \beta_1 X_1 + \dots + \beta_p X_p}_{\text{Controlling vars}} + \underbrace{\alpha_1 Z_1 + \dots + \alpha_q Z_q}_{\text{Explanatory vars}}$

$$H_0: \alpha_1 = \alpha_2 = \alpha_3 = \dots = \alpha_q = 0$$

H_A : At least one α is non-zero

$$\text{logit}(p) = \beta_0 + \beta_1 X_1 + \dots + \beta_p X_p$$

① Max. Likelihood function for M_1, M_2 to get:

$$M_{\text{Full}} \rightarrow \mathcal{L}(\hat{\beta}_0, \hat{\beta}_1, \dots, \hat{\beta}_p, \hat{\alpha}_1, \dots, \hat{\alpha}_q)$$

$$M_{\text{Reduced}} \rightarrow \mathcal{L}(\bar{\beta}_0, \bar{\beta}_1, \dots, \bar{\beta}_p)$$

(notice the hats! diff. est. ∴ diff. models!)

16/ Then use $\mathcal{L}_{Full}$ and $\mathcal{L}_{Reduced}$ to get DEVIANCE!

$$D = (\mathcal{L}_{Reduced} - \mathcal{L}_{Full})$$

This is χ^2 with q d.o.f (The "p"s just 'cancel out' I guess...)

so look for p-value = $P(\chi_q^2 > D)$

Two Things: ① This is approximate! Large samp. size!
② Not a goodness of fit test!

Overall Regression Test: M_{Full} is the same

$$H_0: \beta_1 = \beta_2 = \dots = \beta_r = 0$$

$$M_{Reduced} \text{ is } \text{logit}(p) = \beta_0$$

H_A : At least one β
is non-zero

$$D = (\mathcal{L}_{Reduced} - \mathcal{L}_{Full}), \chi^2 \text{ w/ } p \text{ d.o.f}$$

For both, Reject Null when p is small.

└ Subset: "Adding q things has no effect; the q coeffs are not stat. sig"

└ Overall: "Adding anything has no effect."

Classification: Log Reg is awesome for this!
for $\hat{p}(x_1, \dots, x_p) = \frac{e^{(\hat{\beta}_0 + \hat{\beta}_1 x_1 + \dots + \hat{\beta}_p x_p)}}{1 + e^{(\hat{\beta}_0 + \dots + \hat{\beta}_p x_p)}}$

a classifier $h(x_1, \dots, x_p) = \begin{cases} 1 & \text{if } \hat{p} > 0.5 \\ 0 & \text{if not} \end{cases} \rightarrow \text{can be like } 0.5$

Its error is $L(h) = 1 - \Pr(Y = h(x_1, \dots, x_p))$

Proportion of misclassification is Training Error.

BIASED DOWNWARD! Why? Snooping! Double-dipping!
Use to train ... then use for error analysis!

(Note! Training RISK is biased upward!)

So what to do? Use *Cross-Validation!*
(and no future data X^*)

Split $\swarrow$ testing
 $\searrow$ train!

- LOOCV ($k=n$)
- k-fold
- Random subsampling

Which one to use?

so we

- Random: Expensive, no clear advantage over k-fold
- k-fold: How many folds? 10 is good (LOOCV is extreme...)
Get this higher and you will
 $\downarrow$ BIAS but $\uparrow$ Variance !!

V8
Goodness of Fit $\nleftrightarrow$ Predictive Power!

GOF (Goodness of Fit)

Residual Analysis : Can only do this with
Binary Data with Replication!

Without Replication,

$$P(Y_i | X_i) \sim \text{Bernoulli}(p(X_i)) \quad \text{trial}$$

This is the same thing as $\sim \text{Binomial}(1, p(X_i))$

With Replication, you can see Y_i more than once
for a set of X_{ip} 's

$$\text{So, } P(Y_i | X_{ip}) \sim \text{Binomial}(n_i, p(X_{ip}))$$

So now, we can get
residuals!

trials!
 $n_i > 1$

① Pearson Residuals:
$$r_i = \frac{Y_i - n_i \hat{p}_i}{\sqrt{n_i \hat{p}_i (1 - \hat{p}_i)}}$$

$$\frac{e^{(\hat{p}_i + \hat{p}_i X_i \dots)}}{1 - e^{(\dots)}}$$

L9/ ② Deviance Residuals, d_i

$$d_i = \text{Sign}(Y_i - \hat{Y}_i) \sqrt{\text{some shit} \dots}$$

"saturated" model $\rightarrow$ response is just obs. response!
fitted model.

kinda makes sense in a way...

BOTH have an approx. $SNDist$
if the model assumptions hold.

$\rightarrow$ How to check? - Histogram of Residuals
- Normality (QQ) plot.

If the deviance/residuals fit, Good fit ☺

$\rightarrow$ Also hypothesis test.

H_0 : Model fits the data

H_a : Model does not fit data.

Test stat is $D = \sum d_i^2 \sim \chi^2_{n-p-1}$

as usual reject H_0 if $p\text{-value} < \alpha$ ($\Pr(\chi^2_{df} > D)$)

We actually DONT want to reject H_0 !!

We want high p -values!!

Introduce

In Logistic Reg, to compare models, you are just comparing the **LIKELIHOODS** of each model. That's all.
(that it will explain stuff)

Pred Power!

Subset Testing: $\mathcal{L}(\text{Reduced}) - \mathcal{L}(\text{Full})$

THIS IS ABT
PRED POWER!

GOF Testing: $\mathcal{L}(\text{Saturated}) - \mathcal{L}(\text{Fitted})$

There are separate inferences!

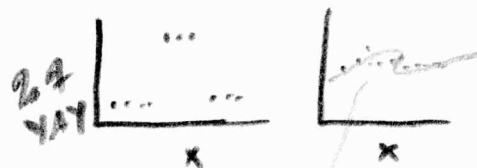
* - Pred Power means its good @ predicting even if one or more assumptions do not hold!

* - GOF $\Rightarrow$ Model Assumption hold

Introduce

- LogR is not necessarily good for any kind of dataset.
Its good for prob., but... use your head :)

- Logit func. must be
linear combo of X's.



↓ GOF
↑ Pred
Superflat
S-curve
↑ GOF
↓ Pred

Curing GOF probs. ① Add pred vars
② Transform vars.

③ Corr. among
responses

③ Outliers...?

④ ~~Not~~ is not appropriate
Binomial dist

⑤ Heterog in Success
that has not been
modeled :)

① and ⑤ lead to OVER DISPERSION

✓✓
Overdispersion: The variability of the prob. estimator is larger than what would be implied by a Binomial Random Var.

So Maybe we another ~~logit~~ s-shape function?

Eg: probit, complementary log-log (c-log-log)

tightest S-shape.
Inverse of CDF of N.D.
"Cumulative"

Fits long tails really well. Use for skewed dist.



But Logit in "canonical" $\rightarrow \hat{\beta}_p$ (param estimator) are fully efficient.

also, can speak/interpret param est. in terms of Log-odds ... cannot with others (probit, c-log-log).

Interlude

Simpson's Paradox: "Reversal of association"

(change in sign of $\hat{\beta}_p$)

when going from marginal $\rightarrow$ cond.

POISSON REGRESSION

Rates, Counts, Shit that fits a Poisson Dist.

~~or an Exp. Dist~~

~~or a family of Exp. Dist~~

But First! GENERALIZED LINEAR MODELS

Pred vars: $n \times p$ matrix X_{np} Response: Y_n

Y_n is a Random Var. from an Exp Dist family

{ Normal (Binomial)
Poisson Gamma

MODEL: "Model some function g as the Expectation of Y as a linear combo of X "

Link function

i.e. $g[E(Y|x_1, \dots, x_p)] = \beta_0 + \beta_1 x_1 + \dots + \beta_p x_p$

or $E(Y|X) = g^{-1}(\beta_0 + \dots + \beta_p x_p)$ Inverse link

Dist	Link	Identity (Regression)	Log ⁻¹
Normal	$g(m) = m$	$m = x^T \beta$	
Poisson	$g(m) = \log(m)$	$m = e^{x^T \beta}$	
Binomial/Bernoulli	$g(m) = \log\left(\frac{m}{1-m}\right)$	$m = \frac{e^{x^T \beta}}{1 + e^{x^T \beta}}$	
Gamma	$g(m) = 1/m$	$m = \frac{1}{x^T \beta}$	

Wrong?
Gamma $\rightarrow$ inverse...?

Intro lude

Why is Normal Dist "Exp."?

Anything is in the "Exp. Family" iff its PDF can be written as:

$$f(y; \theta) = h(y) e^{\underbrace{g(\theta) \cdot T(y)}_{\text{link function}} - B(\theta)}$$

param of dist.

As Poisson: $\log[E(Y|X)] = \beta_0 + \beta_1 x_1 + \dots + \beta_p x_p$

Why is this diff. than standard Reg w/ $\log(Y)$?

Standard: $E[\log(Y)|X]$
 poisson: $\log[E(Y|X)]$ } Flipped!

$E = \mu$ for poisson...

also, Standard: Var: constant
 poisson: $\text{Var}(Y|X) = e^{(\beta_0 + \beta_1 x_1 + \dots + \beta_p x_p)}$
 or $\log(\text{Var}(Y|X)) = \beta_0 + \beta_1 x_1 + \dots + \beta_p x_p$
 NOT constant.

* For rates and counts, if you use StReg with log-transform (and instead of Poisson), this will violate const-var. assumption!

ONLY use if you have to (prefer Poisson) if the counts are large with "variance stabilizer" $\sqrt{\mu + \frac{3}{8}}$

P3 / Model, Estimation, Interpretation

Some Rando Var^Y has Poisson dist with rate λ

$$\text{if } P(Y=y) = \frac{e^{-\lambda} \lambda^y}{y!} \quad (\text{PDF})$$

$$\text{and } E(Y) = V(Y) = \lambda$$

So MODEL: $\log(\lambda_i) = \beta_0 + \beta_1 x_{i1} + \dots + \beta_p x_{ip}$ NO Error term!
NO Variance!
this is "log-rate"

Interpretation: "Log ratio of Rate" change for

Interpret as ratio
of response rates

unit change in β_j
(holding all else fixed)

$$\frac{e^{\beta_0 + \beta_1(x+1)}}{e^{\beta_0 + \beta_1 x}} = e^{\beta_1}$$

Estimation: Use MLE again

(Max log likelihood)

no closed form, numerically determined...

So just like logistic, $\hat{\beta}_j$'s and their errors
are all approx.!

so if $\beta_k = 0.07015$ (Math score) (Award Count)^Y

* unit increase in Math score, the log of the
Expected Award Count inc by 0.07 (all other fixed)

or, The rate ratio of awards would increase
by $e^{0.07015} = 1.072$