

STAT 331



Personal Notes

Marcus Chan

Taught by Peter Bailka
UW CS '25



Chapter 1: Introduction

REGRESSION

💡 In regression modelling, we attempt to explain or account for variation in a response variate (y) by using a model to describe the relationship between y and one or more explanatory variates ($x_1, x_2, \dots$)

SUMMARIES OF THE DATA

💡 A simple LR model involves:

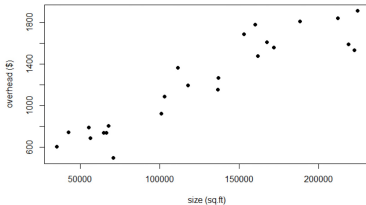
- ① A single explanatory variate: x
- ② A single response variate.

eg Overhead data example:

response (y): claimed overhead (\$)
explanatory (x): office size (sq.ft)

💡 We can summarise the data using a scatter-plot.

Claimed overhead vs office size (n = 24)



💡 To get a numerical summary of the data, we can use the "sample correlation coefficient".

$$r = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum (x_i - \bar{x})^2 \sum (y_i - \bar{y})^2}} = \frac{S_{xy}}{\sqrt{S_{xx} S_{yy}}}$$

Note $-1 \leq r \leq 1$ and that r is unitless.

💡 r tells us the relative strength of the linear relationship.

THE SIMPLE LR MODEL

💡 We can describe the observed behavior of the response with a model that includes both

- ① a "deterministic component" that describes the variation in y accounted for by the functional form of the underlying relationship between y & x ; &

eg with the overhead data, the det. comp. is

$$\mu = \beta_0 + \beta_1 x.$$

where μ = the mean value of y for a given value of x .

- ② an "error term" ϵ that describes the random variation in y not accounted for by the underlying relationship with x .

💡 Putting these together yields the "simple LR" (or SLR) model:

$$y_i = \beta_0 + \beta_1 x_i + \epsilon_i$$

where

- ① β_0 = the "intercept" parameter
- ② β_1 = the "slope" parameter
- ③ i = the index that denotes the observation number.

($x_1, \dots, x_n$ is explanatory data; $y_1, \dots, y_n$ is response data).

* note $\beta_0 + \beta_1 x_i$ is deterministic & ϵ is random.

THE NORMAL SLR MODEL

⚡₁ We typically assume in SLR that

$$\varepsilon_i \stackrel{iid}{\sim} N(0, \sigma^2), \quad i=1, \dots, n$$

for some variance σ^2 .

⚡₂ This yields the normal model

$$y_i = \beta_0 + \beta_1 x_i + \varepsilon_i, \quad \varepsilon_i \stackrel{iid}{\sim} N(0, \sigma^2).$$

⚡₃ Assumptions needed to use this model:

- ① the functional form (ie linear) of the relationship between y & x is correctly specified by the deterministic component of the model;
- ② errors follow a normal distribution;
- ③ errors have a constant variance σ^2 (ie "homoskedasticity"); &
- ④ errors are independent.

LEAST SQUARES ESTIMATION OF MODEL PARAMETERS

⚡₁ Goal: we want to find values of β_0 & β_1 such that for the data

$$\begin{aligned} y_1 &= \beta_0 + \beta_1 x_1 + \varepsilon_1 \\ &\vdots \\ y_n &= \beta_0 + \beta_1 x_n + \varepsilon_n \end{aligned}$$

the sum of squares of the errors $\sum_i \varepsilon_i^2$ is minimized.

⚡₂ The values of β_0 & β_1 obtained by this procedure (denoted $\hat{\beta}_0$ & $\hat{\beta}_1$) are known as the "least squares estimates" of β_0 & β_1 .

⚡₃ We show that

$$\begin{aligned} \hat{\beta}_0 &= \bar{y} - \hat{\beta}_1 \bar{x} \\ \hat{\beta}_1 &= \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2} = \frac{S_{xy}}{S_{xx}}. \end{aligned}$$

Proof. We wish to minimize

$$S(\beta_0, \beta_1) = \sum_{i=1}^n \varepsilon_i^2 = \sum_{i=1}^n [y_i - (\beta_0 + \beta_1 x_i)]^2.$$

See that

$$\frac{\partial S}{\partial \beta_0} = -2 \sum_{i=1}^n [y_i - (\hat{\beta}_0 + \hat{\beta}_1 x_i)]$$

$$\frac{\partial S}{\partial \beta_1} = -2 \sum_{i=1}^n x_i [y_i - (\hat{\beta}_0 + \hat{\beta}_1 x_i)].$$

Since we want to minimize S , we can solve

$$\begin{cases} \frac{\partial S}{\partial \beta_0} = 0 \\ \frac{\partial S}{\partial \beta_1} = 0. \end{cases}$$

The resultant solutions for β_0 & β_1 are the desired values as required. □

⚡₄ In R, we can get these values via

> data.slr.lm <- lm(response ~ explanatory).

FITTED MODEL

⚡₁ For the SLR model, the fitted model is

$$\hat{\mu} = \hat{\beta}_0 + \hat{\beta}_1 x,$$

where $\hat{\mu}$ is the estimated mean value of y given a value of x .

FITTED RESIDUALS

⚡₁ The "fitted residual" of the i th observation, e_i , is defined as

$$e_i = y_i - \hat{\mu}_i = y_i - (\hat{\beta}_0 + \hat{\beta}_1 x_i).$$

* ε_i is a random variable in which we impose assumptions;
 e_i is the difference between the observed response & estimated mean response.

⚡₂ If we take the partial derivative wrt each parameter and set = 0 in our least squares procedure, we get that

$$\begin{aligned} \sum e_i &= 0 \\ \sum x_i e_i &= 0. \end{aligned}$$

⚡₃ These constraints allow us to calculate the remaining 2 residuals from $n-2$ observations;

so we say the fitted model is associated with $n-2$ degrees of freedom.

LEAST SQUARES ESTIMATE OF $\sigma^2: \hat{\sigma}^2$

💡₁ In the normal model, we assume

$$e_i \stackrel{iid}{\sim} N(0, \sigma^2).$$

💡₂ In any least squares regression model, we estimate σ^2 by dividing the sum of squares of the residuals by the degrees of freedom.

💡₃ In particular, this means our estimate for σ^2 is

$$\hat{\sigma}^2 = \frac{\sum_{i=1}^n e_i^2}{n-2} = \frac{\sum_{i=1}^n (y_i - \hat{\mu}_i)^2}{n-2}.$$

* note $E[\hat{\sigma}^2] = \sigma^2$ (ie $\hat{\sigma}^2$ is unbiased).

RESIDUAL STANDARD ERROR: $\hat{\sigma}$

💡 The "residual standard error" is

$$\hat{\sigma} = \sqrt{\hat{\sigma}^2} = \sqrt{\frac{\sum e_i^2}{n-2}}$$

💡₂ $\hat{\sigma}$ can be interpreted as the estimated std dev of the errors & measures the random variation of the response given a value for x.

💡₃ The smaller $\hat{\sigma}$ is, the more the variation in y is "explained" by x, and so the better fit the model is.

💡₄ $\hat{\sigma}$ is part of the summary R output for the fitted model:

```
> summary(audit.lm)
Call:
lm(formula = overhead ~ size)

Residuals:
    Min       1Q   Median       3Q      Max
-36639 -12874  -1997   8642  56686

Coefficients:
            Estimate Std. Error t value Pr(>|t|)
(Intercept) -27877.06   14172.00  -1.967   0.0619 .
size         126.33     10.88   11.610 7.47e-11 ***
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 23480 on 22 degrees of freedom
Multiple R-squared:  0.8597,    Adjusted R-squared:  0.8533
F-statistic: 134.8 on 1 and 22 DF,  p-value: 7.472e-11
```

INTERPRETATION OF PARAMETER ESTIMATES

💡₁ We may interpret $\hat{\beta}_1$ as the estimated mean change in the response y associated with a change of one unit in x.

💡₂ We may interpret $\hat{\beta}_0$ as the estimated mean value of y at x=0 only if x=0 is a relevant value and is in the range of values we used to fit the model.

* never extrapolate to values of x outside the range used to fit the model.

💡₃ Lastly, we can interpret $\hat{\sigma}$ as a measure of the variability of the response about the fitted line.

INFERENCE FOR β_1

💡₁ To investigate whether there is a linear relationship between y & x in the population, we can test the hypothesis " $\beta_1 = 0$."

💡₂ We can then either use confidence intervals or hypothesis tests to test this.

💡₃ To do this, we need the least squares estimator of β_1 :

$$\hat{\beta}_1 = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2}$$

DISTRIBUTION OF $\hat{\beta}_1$

💡 We can show for the SLR model that

$$\hat{\beta}_1 \sim N(\beta_1, \frac{\sigma^2}{S_{xx}})$$

Proof. First, note

$$\begin{aligned}\hat{\beta}_1 &= \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2} \\ &= \frac{\sum (x_i - \bar{x})y_i - \bar{y} \sum (x_i - \bar{x})}{\sum (x_i - \bar{x})^2} \\ &= \frac{\sum (x_i - \bar{x})y_i}{\sum (x_i - \bar{x})^2} \quad \because \sum (x_i - \bar{x}) = 0 \\ &= \sum c_i y_i, \quad c_i = \frac{x_i - \bar{x}}{\sum (x_i - \bar{x})^2}\end{aligned}$$

Then, for the SLR model, $\epsilon_i \stackrel{iid}{\sim} N(0, \sigma^2)$.

Since $y_i = \beta_0 + \beta_1 x_i + \epsilon_i$, thus

$y_i \sim N(\beta_0 + \beta_1 x_i, \sigma^2)$ & y_i are ind.

and so

$\hat{\beta}_1 = \sum c_i y_i \sim \text{Normal}$ (by properties of normal).

Then

$$\begin{aligned}E(\hat{\beta}_1) &= E(\sum c_i y_i) = \sum c_i E(y_i) \\ &= \sum \frac{x_i - \bar{x}}{\sum (x_i - \bar{x})^2} \cdot (\beta_0 + \beta_1 x_i) \\ &= \frac{\beta_0 \sum (x_i - \bar{x}) + \beta_1 \sum x_i (x_i - \bar{x})}{\sum (x_i - \bar{x})^2} \\ &= \frac{\beta_1 \sum x_i (x_i - \bar{x})}{\sum (x_i - \bar{x})^2} \\ &= \frac{\beta_1 \sum x_i (x_i - \bar{x}) - \beta_1 \bar{x} \sum (x_i - \bar{x})}{\sum (x_i - \bar{x})^2} \\ &= \frac{\beta_1 \sum (x_i - \bar{x})^2}{\sum (x_i - \bar{x})^2} = \beta_1\end{aligned}$$

Similarly,

$$\begin{aligned}\text{Var}(\hat{\beta}_1) &= \text{Var}(\sum c_i y_i) \\ &= \sum c_i^2 \text{Var}(y_i) \quad \because y_i \text{'s ind.} \\ &= \sum \frac{(x_i - \bar{x})^2}{(\sum (x_i - \bar{x})^2)^2} \cdot \sigma^2 \\ &= \frac{\sigma^2}{\sum (x_i - \bar{x})^2} = \frac{\sigma^2}{S_{xx}}\end{aligned}$$

Hence $\beta_1 \sim N(\beta_1, \frac{\sigma^2}{S_{xx}})$ as required.

💡 It follows that

$$\frac{\hat{\beta}_1 - \beta_1}{SE(\hat{\beta}_1)} = \frac{\hat{\beta}_1 - \beta_1}{(\hat{\sigma} / \sqrt{S_{xx}})} \sim t_{n-2}$$

(from STAT231/330 result).

This can be used to get t-based CIs & hypothesis tests for β_1 .

DISTRIBUTION OF $\hat{\beta}_0$

💡 Similarly, we can show in a SLR model,

$$\begin{aligned}\hat{\beta}_0 &\sim N(\beta_0, \sigma^2(\frac{1}{n} + \frac{\bar{x}^2}{S_{xx}})) \\ \frac{\hat{\beta}_0 - \beta_0}{SE(\hat{\beta}_0)} &= \frac{\hat{\beta}_0 - \beta_0}{\hat{\sigma} \sqrt{\frac{1}{n} + \frac{\bar{x}^2}{S_{xx}}}} \sim t_{n-2}\end{aligned}$$

CI FOR β_1

💡 A $(1-\alpha)100\%$ confidence interval for β_1 is

$$\hat{\beta}_1 \pm t_{n-2, 1-\alpha/2} SE(\hat{\beta}_1), \quad SE(\hat{\beta}_1) = \frac{\hat{\sigma}}{\sqrt{S_{xx}}}$$

- $t_{n-2, 1-\alpha/2}$:= the critical value from a t_{n-2} distribution corresponding to a confidence level of $(1-\alpha)100\%$.

💡 " $t_{n-2, 1-\alpha/2} SE(\hat{\beta}_1)$ " is called the "margin of error" of the interval.

💡 We can use R to calculate this:

> summary(data.lm)

```
Estimate Std. Error t value Pr(>|t|)
(Intercept) -27877.06 14172.00 -1.967 0.0619
size 126.33 10.88 11.610 7.47e-11
```

> $t \leftarrow qt(1 - \frac{\alpha}{2}, n-2)$

↳ The CI is then $126.33 - t(10.88)$, $126.33 + t(10.88)$.

💡 We may interpret the CI as that we are $(1-\alpha)100\%$ confident that for every additional increase of a unit of x , the mean increase of y is between (start of CI) & (end of CI).

💡 If " $\beta_1 = 0$ " is not in the interval, then we say there is a significant relationship between x & y (at the $(1-\alpha)100\%$ confidence level).

💡 Hypothesis test for β_1 :

- ① $H_0: \beta_1 = 0$; $H_A: \beta_1 \neq 0$
- ② (Assuming H_0) our test statistic is

$$t = \frac{\hat{\beta}_1 - \beta_1}{SE(\hat{\beta}_1)} = \frac{\hat{\beta}_1}{SE(\hat{\beta}_1)}$$

- ③ p-value is $p = 2P(T > t)$, $T \sim t_{n-2}$
→ In R:
> $p \leftarrow 2 * (1 - pt(t, n-2))$
- ④ Check if $p < 0.05$; if yes, reject H_0 .

Chapter 2:

Multiple Regression

MULTIPLE REGRESSION MODEL

💡 If we expand the SLR model to p explanatory variables, we obtain the multiple linear regression model:

$$y_i = \beta_0 + \beta_1 x_{i1} + \dots + \beta_p x_{ip} + \varepsilon_i, \quad i=1, \dots, n$$

💡 This can be expressed as

$$y = X\beta + \varepsilon$$

where $y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} \in \mathbb{R}^n$, $X = \begin{pmatrix} 1 & x_{11} & \dots & x_{1p} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & x_{n1} & \dots & x_{np} \end{pmatrix} \in \mathbb{R}^{n \times (p+1)}$,

$$\beta = \begin{pmatrix} \beta_0 \\ \vdots \\ \beta_p \end{pmatrix} \in \mathbb{R}^{p+1} \quad \& \quad \varepsilon = \begin{pmatrix} \varepsilon_1 \\ \vdots \\ \varepsilon_n \end{pmatrix} \in \mathbb{R}^n.$$

NORMAL MODEL

💡 For the normal model, where we assume $\varepsilon_i \stackrel{iid}{\sim} N(0, \sigma^2)$, we write

$$y = X\beta + \varepsilon, \quad \varepsilon \sim \text{MVN}(0, \sigma^2 I),$$

where $\text{Var}(\varepsilon) = \sigma^2 I$ is the covariance matrix of the error random vector ε .

LEAST SQUARES ESTIMATION OF β

💡 We wish to minimize

$$S(\beta_0, \dots, \beta_p) = \sum_{i=1}^n \varepsilon_i^2 = \sum [y_i - (\beta_0 + \beta_1 x_{i1} + \dots + \beta_p x_{ip})]^2$$

over $\beta_0, \dots, \beta_p$.

💡 Taking partial derivatives and setting to 0:

$$\frac{\partial S}{\partial \beta_0} = -2 \sum (y_i - (\hat{\beta}_0 + \hat{\beta}_1 x_{i1} + \dots + \hat{\beta}_p x_{ip})) = 0$$

$$\frac{\partial S}{\partial \beta_1} = -2 \sum x_{i1} (y_i - (\hat{\beta}_0 + \hat{\beta}_1 x_{i1} + \dots + \hat{\beta}_p x_{ip})) = 0$$

$\vdots$

$$\frac{\partial S}{\partial \beta_p} = -2 \sum x_{ip} (y_i - (\hat{\beta}_0 + \hat{\beta}_1 x_{i1} + \dots + \hat{\beta}_p x_{ip})) = 0$$

💡 This yields the normal equations

$$n(\hat{\beta}_0) + \hat{\beta}_1 \sum x_{i1} + \dots + \hat{\beta}_p \sum x_{ip} = \sum y_i$$

$$\hat{\beta}_0 \sum x_{i1} + \hat{\beta}_1 \sum x_{i1}^2 + \dots + \hat{\beta}_p \sum x_{i1} x_{ip} = \sum x_{i1} y_i$$

$\vdots$

$$\hat{\beta}_0 \sum x_{ip} + \hat{\beta}_1 \sum x_{i1} x_{ip} + \dots + \hat{\beta}_p \sum x_{ip}^2 = \sum x_{ip} y_i$$

💡 We can write this as

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

and so

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

- note this needs $X^T X$ to be invertible;
ie full rank / all columns are linearly independent.

💡 Note:

① The fitted line is given by

$$\hat{\mu} = x^T \hat{\beta}.$$

② The vector of fitted values is

$$\hat{\mu} = X \hat{\beta}.$$

③ The residual vector is

$$e = y - \hat{\mu}$$

* sum of squares of residuals is $\sum e_i^2 = e^T e$.

THE HAT MATRIX: H

💡 We can express $\hat{\mu}$ by

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

where

$$H = X(X^T X)^{-1} X^T y$$

is the "hat" matrix which maps the vector of response variables to the vector of fitted values.

💡 Note that

① H is symmetric (ie $H^T = H$); &

② H is idempotent (ie $H^2 = H$).

💡 We can express our residual vector e as

$$e = y - \hat{\mu} = y - Hy = (I - H)y$$

LEAST SQUARES ESTIMATION OF σ^2

💡 The least squares estimate of σ^2 for a p explanatory variable multiple regression model with (p+1) parameters is

$$\hat{\sigma}^2 = \frac{\sum_{i=1}^n e_i^2}{n - (p+1)}$$

where $df = n - (p+1)$.

RESIDUAL STANDARD ERROR

💡 The residual standard error is thus

$$\hat{\sigma} = \sqrt{\frac{\sum e_i^2}{n - (p+1)}}$$

MLE FOR β

💡 The MLE for β is equivalent to the least squares estimate; ie the likelihood function

$$\begin{aligned} L(\beta_0, \dots, \beta_n | y_1, \dots, y_n) &= \prod_{i=1}^n f(y_i) \\ &= \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(y_i - \mu_i)^2}{2\sigma^2}}, \quad \mu_i = \beta_0 + \sum_j \beta_j x_{ij} \\ &= (2\pi\sigma^2)^{-n/2} \exp\left(-\frac{\sum (y_i - \mu_i)^2}{2\sigma^2}\right) \end{aligned}$$

or equivalently the log likelihood function

$$\ell(\beta_0, \dots, \beta_n | y_1, \dots, y_n) = c - \frac{\sum_i (y_i - (\beta_0 + \beta_1 x_{i1} + \dots + \beta_p x_{ip}))^2}{2\sigma^2}$$

is maximized at $\beta = (\beta_0, \dots, \beta_p)$ that minimizes

$$\sum e_i^2 = \sum (y_i - (\beta_0 + \beta_1 x_{i1} + \dots + \beta_p x_{ip}))^2$$

GAUSS-MARKOV THEOREM & BLUE

💡 The least squares estimator

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

is the "best linear unbiased estimator"

(BLUE) of β .

💡 More formally, if we consider the model given by

$$y = X\beta + \epsilon, \quad E(\epsilon) = 0, \quad \text{Var}(\epsilon) = \sigma^2 I$$

then amongst all unbiased linear estimators

$\hat{\beta}^* = M^* y$, the least squares estimator

$\hat{\beta} = My$ has the "smallest" variance;

ie

$$\text{Var}(\hat{\beta}^*) = \text{Var}(\hat{\beta}) + \sigma^2 (M^* - M)(M^* - M)^T$$

where $(M^* - M)(M^* - M)^T$ is positive semidefinite.

- A is "positive definite" if $a^T A a > 0$ for any vector a.

DISTRIBUTION OF $\hat{\beta}$

💡 We show that

$$\hat{\beta} \sim \text{MVN}(\beta, \sigma^2(X^T X)^{-1})$$

where MVN is the multivariate normal distribution.

Proof. First, we have

$$Y = X\beta + \epsilon, \quad \epsilon \sim \text{MVN}(0, \sigma^2 I).$$

Thus, by properties of MVN,

$$Y \sim \text{MVN}(X\beta, \sigma^2 I).$$

Hence $\hat{\beta} = (X^T X)^{-1} X^T Y$ also follows a MVN distribution.

Next, see that

$$\begin{aligned} E(\hat{\beta}) &= E((X^T X)^{-1} X^T Y) \\ &= (X^T X)^{-1} X^T E(Y) \\ &= (X^T X)^{-1} X^T (X\beta) \\ &= \beta. \end{aligned}$$

Then

$$\begin{aligned} \text{Var}(\hat{\beta}) &= \text{Var}((X^T X)^{-1} X^T Y) \\ &= (X^T X)^{-1} X^T \text{Var}(Y) [(X^T X)^{-1} X^T]^T \\ &\quad (\text{Var}(AY) = A \text{Var}(Y) A^T) \\ &= \sigma^2 (X^T X)^{-1} X^T [(X^T X)^{-1} X^T]^T \\ &= \sigma^2 (X^T X)^{-1} X^T X (X^T X)^{-1} \\ &= \sigma^2 (X^T X)^{-1}, \end{aligned}$$

which gives us the desired result.

💡 The marginal distribution of $\hat{\beta}_j$ is thus

$$\hat{\beta}_j \sim N(\beta_j, \sigma^2 (X^T X)^{-1}_{jj}) \quad \forall j=0, \dots, p$$

💡 We also have that

$$\frac{\hat{\beta}_j - \beta_j}{\text{SE}(\hat{\beta}_j)} \sim t_{n-(p+1)}, \quad \text{SE}(\hat{\beta}_j) = \hat{\sigma} \sqrt{(X^T X)^{-1}_{jj}}$$

💡 Also note that

$$\text{Cov}(\hat{\beta}_i, \hat{\beta}_j) = \sigma^2 (X^T X)^{-1}_{ij}.$$

INTERPRETATION OF $\hat{\beta}_j$

💡 $\hat{\beta}_j$ is the estimated mean change in the response associated with a change of one unit of x_j whilst holding all other variables constant.

CI'S FOR β_j

💡 A $(1-\alpha) 100\%$ CI for β_j is

$$\hat{\beta}_j \pm t_{n-(p+1), 1-\alpha/2} \text{SE}(\hat{\beta}_j)$$

💡 If $\beta_j = 0$ is not in the CI, then there is a significant linear relationship between y & x_j .

HYPOTHESIS TESTS FOR β_j

💡 Hypothesis test for β_j :

$$① H_0: \beta_j = 0; \quad H_A: \beta_j \neq 0$$

$$② t = \frac{\hat{\beta}_j - \beta_j}{\text{SE}(\hat{\beta}_j)} = \frac{\hat{\beta}_j}{\text{SE}(\hat{\beta}_j)}$$

$$③ p\text{-value} = 2P(T > t), \quad T \sim t_{n-(p+1)}$$

$$④ \text{Reject } H_0 \text{ if } p < 0.05.$$

MULTI-COLLINEARITY

💡₁ We say 2 or more explanatory variables exhibit "multicollinearity" if there exist strong linear relationships between them.

💡₂ This

- ① increases the variances (and thus std errors) of the associated parameter estimators; &
- ② leads to wide/imprecise CIs & inaccurate conclusions from hypothesis tests.

VARIANCE INFLATION FACTOR / VIF

💡₁ The "variance inflation factor" is a measure of multicollinearity associated with some explanatory variable x_j .

💡₂ How to calculate VIF for x_j :

- ① Regress x_j onto all other x 's; ie fit models for x_j against each other x ;

- ② Then

$$\text{VIF}_j = \frac{1}{1 - R_j^2},$$

where R_j^2 is the coefficient of determination of the model fit with x_j as the response.

- in R_j^2 is the "multiple R squared" parameter.

💡₃ Generally, we remove x_j from the model if $\text{VIF}_j > 10 \Leftrightarrow R_j^2 > .90$.

💡₄ In R , we can do

$$> \text{lm}(x \sim \underbrace{x_1 + x_2 + \dots + x_n}_{\substack{\uparrow \\ \text{variable we} \\ \text{are testing}}} \text{ other exp. variables})$$

and check the multiple R-squared value.

CI FOR μ_{new}

💡₁ Idea: We may want to use our fitted model to estimate the mean response of a new unit in the population.

💡₂ In particular,

$$\hat{\mu}_{\text{new}} = x_{\text{new}}^T \hat{\beta}.$$

💡₃ Then, we show that

$$\hat{\mu}_{\text{new}} \sim N(x_{\text{new}}^T \beta, \sigma^2 x_{\text{new}}^T (X^T X)^{-1} x_{\text{new}})$$

Proof. Recall

$$\hat{\beta} \sim \text{MVN}(\beta, \sigma^2 (X^T X)^{-1}).$$

Then $\hat{\mu}_{\text{new}} = x_{\text{new}}^T \hat{\beta}$ must also follow a normal distribution.
See that

$$\begin{aligned} E(\hat{\mu}_{\text{new}}) &= E(x_{\text{new}}^T \hat{\beta}) \\ &= x_{\text{new}}^T E(\hat{\beta}) \\ &= x_{\text{new}}^T \beta \end{aligned}$$

Then

$$\begin{aligned} \text{Var}(\hat{\mu}_{\text{new}}) &= \text{Var}(x_{\text{new}}^T \hat{\beta}) \\ &= x_{\text{new}}^T \text{Var}(\hat{\beta}) x_{\text{new}} \\ &= \sigma^2 x_{\text{new}}^T (X^T X)^{-1} x_{\text{new}}. \end{aligned}$$

This gives the result, so we're done. $\square$

💡₄ Thus, a $(1-\alpha)\%$ confidence interval for μ_{new} is

$$\hat{\mu}_{\text{new}} \pm t_{n-(p+1), 1-\frac{\alpha}{2}} \hat{\sigma} \sqrt{x_{\text{new}}^T (X^T X)^{-1} x_{\text{new}}}$$

PREDICTION INTERVAL FOR

y_{new}

💡 Idea: We may also wish to use our fitted model to predict the value of the response of a new unit of the population.

💡 Then, note the variance of $\hat{y}_{\text{new}}$ is composed of 2 sources of variation:

- ① the variation associated with the parameter estimators; &
- ② the variance σ^2 associated with a random response.

💡 Thus our total variation is

$$\sigma^2 + \sigma^2 x_{\text{new}}^T (X^T X)^{-1} x_{\text{new}}.$$

💡 Hence, a $(1-\alpha)100\%$ prediction interval for y_{new} is

$$\hat{y}_{\text{new}} \pm t_{n-(p+1), 1-\frac{\alpha}{2}} \hat{\sigma} \sqrt{1 + x_{\text{new}}^T (X^T X)^{-1} x_{\text{new}}}$$

CONFIDENCE & PREDICTION BANDS FOR THE SLR MODEL

💡 For the SLR model,

$$x_{\text{new}}^T (X^T X)^{-1} x_{\text{new}} = \frac{1}{n} + \frac{(x_{\text{new}} - \bar{x})^2}{S_{xx}}$$

💡 Thus, the closer x_{new} is to $\bar{x}$, the narrower the CIs & PIs.

💡 In general, the closer $\{x_1, \dots, x_p\}$ is to $\{\bar{x}_1, \dots, \bar{x}_p\}$ in the multiple regression model, the narrower the interval.

MODELLING CATEGORICAL EXPLANATORY VARIABLES

💡 We can code categorical explanatory variables using indicator variables that take on values of 0 or 1.

$$\text{eg } x_1 = \mathbb{I}[A=a_1], \quad x_2 = \mathbb{I}[A=a_2]$$

$\mathbb{I}$ = the indicator function.

💡 For a variable with L category levels, we need $L-1$ indicator variables.

💡 In particular, if X is a categorical variable that has L distinct values

$a_1, \dots, a_L$, we can use the model

$$\hat{\mu} = \hat{\beta}_0 + \hat{\beta}_1 x_1 + \dots + \hat{\beta}_{L-1} x_{L-1},$$

where $x_i = \mathbb{I}[X=a_i]$.

$-(x_1, \dots, x_{L-1}) = (0, \dots, 0)$ corresponds with

$X = a_L$.

💡 Then, each $\hat{\beta}_i$ corresponds to the difference in the estimated mean value of the response where $X=a_i$ relative to where $X \neq a_i$.

INFERENCE FOR PARAMETERS ASSOCIATED WITH INDICATOR VARIABLES

💡 To test whether there is a difference in $\hat{\mu}$ between data where $x_i=1$ vs. $x_i=0$, we can use the following hypothesis test:

- ① $H_0: \beta_i = 0; H_A: \beta_i \neq 0$
- ② We get t, p from the 'summary' output from the model fit (ie lm); &
- ③ If $p < 0.05$, we reject H_0 .

💡 To test whether there is a difference in $\hat{\mu}$ between data where $x_i=1$ vs. $x_j=1$, we can use the following hypothesis test:

- ① $H_0: \beta_i - \beta_j = 0, H_A: \beta_i - \beta_j \neq 0$
- ② We can derive

$$\hat{\beta}_i - \hat{\beta}_j \sim N(\beta_i - \beta_j, \sigma^2((X^T X)^{-1}_{ii} + (X^T X)^{-1}_{jj} - 2(X^T X)^{-1}_{ij}))$$

③ Thus we get

$$t = \frac{(\hat{\beta}_i - \hat{\beta}_j)}{\text{SE}(\hat{\beta}_i - \hat{\beta}_j)}, \quad p = 2P(T > t), \quad T \sim t_{n-(p+1)}$$

ORTHOGONAL X MATRIX DESIGNS

Q1 We may wish to model categorical variables in a way that creates an orthogonal X matrix, thus producing independent parameter estimators.

eg Suppose we had 3 categorical variables X_1, X_2, X_3 , each with 2 values 0 & 1.

We can define

$$X_i = \begin{cases} 1, & X_i = 1 \\ -1, & X_i = 0 \end{cases}$$

This yields the "X" matrix

$$X = \begin{pmatrix} | & | & | & | \\ -1 & 1 & 1 & 1 \\ 1 & 1 & -1 & 1 \\ 1 & -1 & 1 & 1 \\ -1 & -1 & -1 & 1 \end{pmatrix}$$

Note the columns of X are orthogonal, and in particular

$$(X^T X) = 8I_4, \quad (X^T X)^{-1} = \frac{1}{8}I_4, \text{ where}$$

I_4 is the 4x4 identity matrix.

Hence $\sigma^2 (X^T X)^{-1} = \text{Var}(\hat{\beta})$ is diagonal, indicating independent parameter estimators for the normal model.

Q2 In particular, when comparing one level to another in X for a given factor, it corresponds to an estimated mean change in the response of $2\hat{\beta}_j$.

- since parameter estimators are independent, they are unaffected by inclusion/exclusion of other variables

- so we do not have to account for other variables when interpreting the parameter estimates associated with a given factor

ANALYSIS OF VARIANCE (ANOVA)

Q1 We can express the total sum of squares of the observations y_i as

$$SS(\text{Tot}) = \sum (y_i - \bar{y})^2 = \underbrace{\sum (\hat{\mu}_i - \bar{y})^2}_{SS(\text{Reg})} + \underbrace{\sum (y_i - \hat{\mu}_i)^2}_{SS(\text{Res})}$$

Q2 In particular,

① The "regression sum of squares"

$SS(\text{Reg})$ is the variation explained by the model; &

② The "residual sum of squares"

$SS(\text{Res})$ is the variation in the response left unexplained by the model.

Q3 In ANOVA methods of inference, we draw conclusions about the relative fit of models by comparing $SS(\text{Reg})$ & $SS(\text{Res})$.

Q4 The greater $SS(\text{Reg})$ is compared to $SS(\text{Res})$, the better the model fit.

COEFFICIENT OF DETERMINATION / MULTIPLE R-SQUARED: R^2

Q1 The "coefficient of determination" is

$$R^2 = 1 - \frac{SS(\text{Res})}{SS(\text{Tot})}$$

Q2 R^2 measures the proportion of the variation in the response explained by the model.

F-TEST FOR MODEL PARAMETERS:

$$\beta_1 = \dots = \beta_p = 0$$

Q1 Idea: To test if a relationship exists between the response at least one of the explanatory variables, we can use a

F-test.

Method:

① $H_0: \beta_1 = \dots = \beta_p = 0; H_A: \exists j \text{ s.t. } \beta_j \neq 0$

② Test statistic:

$$F = \frac{SS(\text{Reg})/p}{SS(\text{Res})/(n-p-1)} = \frac{MS(\text{Reg})}{MS(\text{Res})}$$

③ Under H_0 , F has a F distribution with $p, n-p-1$ degrees of freedom.

④ The F-test statistic & p-value are provided on the last line of the 'summary' output for lm .

⑤ We reject H_0 if $p < 0.05$.

ANOVA TABLE

⚡ We summarise the test $H_0: \beta_1 = \dots = \beta_p = 0$ in an ANOVA table:

Source	df	SS	MS	F	p-value
regression	p	SS(Reg)	$\frac{SS(Reg)}{p}$	$\frac{MS(Reg)}{MS(Res)}$	$P(F_{p, n-(p+1)} > F)$
residual	$n-(p+1)$	SS(Res)	$\frac{SS(Res)}{n-(p+1)}$		
total	$n-1$	SS(Tot)			

⚡ We can obtain SS(Res) from $\hat{\sigma}^2$:

$$\hat{\sigma}^2 = \sqrt{\frac{\sum e_i^2}{n-(p+1)}} = \sqrt{\frac{SS(Res)}{n-(p+1)}}$$

$$\Rightarrow SS(Res) = (n-(p+1)) \hat{\sigma}^2$$

ADDITIONAL SUM OF SQUARES

⚡ Consider the "full" model of

$$Y = \beta_0 + \beta_1 x_1 + \dots + \beta_p x_p + \epsilon, \quad \epsilon \sim N(0, \sigma^2)$$

and the "reduced" model that reflects the restrictions imposed by $\beta_1 = \dots = \beta_p = 0$, $k \leq p$, ie

$$Y = \beta_0 + \beta_{k+1} x_{k+1} + \dots + \beta_p x_p + \epsilon.$$

ie "After accounting for $\beta_1, \dots, \beta_k$, does ≥ 1 of $\beta_{k+1}, \dots, \beta_p$ account for significant variation in y ?"

⚡ To determine which model is better, we can examine

$$SS(Reg)_{full} - SS(Reg)_{red}$$

or $SS(Res)_{red} - SS(Res)_{full}$

which is the difference in the variation explained by the full & reduced models.

⚡ We can test $H_0: \beta_1 = \dots = \beta_k = 0$ with the F test statistic

$$F = \frac{(SS(Res)_{red} - SS(Res)_{full}) / df_{red} - df_{full}}{SS(Res)_{full} / df_{full}}$$

where under H_0 , $F \sim F_{df_{red}-df_{full}, df_{full}}$.

⚡ Our p-value is $P(F \geq F_{test})$, $F \sim F_{df_{red}-df_{full}, df_{full}}$.

⚡ In R, we can use the `anova` function to do this:

```
> anova(audit.red.lm, audit.full.lm)
Analysis of Variance Table
Model 1: overhead ~ col + clients
Model 2: overhead ~ size + age + col + clients
  Res.Df RSS      Df Sum of Sq  F      Pr(>F)
1      21 4954374034
2      19 3901347198  2    1.053e+09  2.5642 0.1033
```

⚡ Note that our previous ANOVA method to test $H_0: \beta_1 = \dots = \beta_p = 0$ is just an additional sum of squares test where our reduced model is

$$Y = \beta_0 + \epsilon, \quad \epsilon \sim N(0, \sigma^2).$$

⚡ Similarly, we can also test $H_0: \beta_j = 0$ using the additional sum of squares test statistic

$$F = \frac{SS(Res)_{red} - SS(Res)_{full}}{SS(Res)_{full} / df_{full}}$$

where our reduced model is just the full model with x_j omitted & under H_0 ,

$$F \sim F_{1, p}.$$

(Note $P(1t_{1,v} > 1t_1) = P(F_{1,v} > t^2)$.)

ADDITIONAL SUM OF SQUARES & CATEGORICAL VARIABLES

💡 Suppose we have a model

$$Y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \epsilon$$

where $x_i = \mathbb{I}[A = a_i]$, where A is a categorical variable.

💡 To test $H_0: \beta_1 = \beta_2 = 0$, we can use the reduced model under $H_0: \beta_1 = \beta_2 = \beta^*$; ie

$$Y = \beta_0 + \beta^* (x_1 + x_2) + \epsilon \\ = \beta_0 + \beta^* x^* + \epsilon$$

where $x^* = \mathbb{I}[A = a_1 \text{ or } A = a_2]$.

💡 We can then use the additional sum of squares statistic & proceed as before.

GENERAL LINEAR HYPOTHESIS

💡 The "general linear hypothesis" is in the form

$$H_0: A\beta = 0$$

where $A \in \mathbb{R}^{L \times (p+1)}$ describes the L linear constraints on the full model as described by H_0 .

💡 In particular, we can use the additional sum of squares test statistic to test H_0 under H_0 ,

$$F = \frac{(SS(Res)_{red} - SS(Res)_{full}) / (df_{red} - df_{full})}{SS(Res)_{full} / df_{full}} \sim F_{df_{red} - df_{full}, df_{full}}$$

eg $H_0: \beta_1 = \beta_2 = \beta_3 = \beta_4 = 0$ can be expressed as

$$H_0: \underbrace{\begin{pmatrix} 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}}_A \underbrace{\begin{pmatrix} \beta_0 \\ \beta_1 \\ \beta_2 \\ \beta_3 \\ \beta_4 \end{pmatrix}}_{\beta} = \underbrace{\begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}}_0$$

ASSESSING MODEL ADEQUACY / RESIDUAL ANALYSIS

💡 Idea: We can assess the "model adequacy" (ie the validity of the model assumptions) by examining the fitted residuals $e = y - \hat{\mu}$.

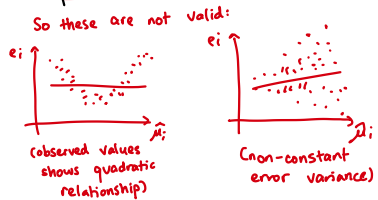
RESIDUAL PLOTS

💡 We can plot the residuals e_i against the fitted values $\hat{\mu}_i$.



💡 If the model assumptions hold, then e_i and $\hat{\mu}_i$ should be uncorrelated.

💡 Thus, if the model assumptions hold, we should see no observable pattern in the plot.



QQ PLOTS

💡 In a QQ plot, we plot the ordered residuals $e_{(i)}$ against the expected ordered values $E(Z_{(i)})$, $Z_i \sim N(0, 1)$.

💡 We use QQ plots to assess the assumption of normal errors.

💡 In particular, if the errors (and hence residuals) are from a normal distribution, then $e_{(i)}$ should be proportional to $E(Z_{(i)})$, so the plot should exhibit a straight line relationship.

VARIANCE STABILIZING TRANSFORMATIONS

💡 "Variance stabilizing transformations" are transformations to the response and possibly some of the explanatory variables to improve model adequacy / validity of the model assumptions.

$$\text{eg } \log(y), y^{1/2}, y^{-1}$$

💡 These are useful when we can write

$$\sigma^2 = f(\mu).$$

PROPERTIES OF THE RESIDUALS

💡 Recall we can express

$$e = (I - H)e$$

where H is the hat matrix ($H = X(X^T X)^{-1} X^T$).

💡 Since $e \sim N(0, \sigma^2 I)$, it follows that

$$e \sim N(0, \sigma^2(I - H)) \\ \Rightarrow e_i \sim N(0, \sigma^2(1 - h_{ii}))$$

where h_{ii} is the i^{th} diagonal element of H .

💡 Note

$$① \text{Var}(e_i) = \sigma^2(1 - h_{ii})$$

- residuals have non-constant variance

$$② \text{Cov}(e_i, e_j) = -\sigma^2 h_{jk}, \quad j \neq i$$

- residuals are not independent

- result of constraint $\sum e_i = 0$ in least squares estimation

STUDENTIZED RESIDUALS

💡 For a random variable X , the "studentized" version of X is

$$Y = \frac{X - \mu}{\hat{\sigma}}$$

where μ is the mean and $\hat{\sigma}$ is an estimate of the standard deviation.

💡 A "studentized residual" associated with observation i is

$$d_i = \frac{e_i}{\hat{\sigma} \sqrt{1 - h_{ii}}}$$

which is $\sim N(0, 1)$ for large n .

EXTREME RESPONSE VALUES / OUTLIERS

💡 An "outlier" is an observation for which $e_i = y_i - \hat{\mu}_i$ is extreme relative to the other residuals.

💡 In particular, we can consider an observation an outlier if

$$|d_i| > 2.5.$$

💡 Causes of outliers:

① Typos / misrecording of data;

② Values of associated potential explanatory variables not included in the model;

③ Random variability.

LEVERAGE

💡 "Leverage" is a measure used to identify those observations whose set of explanatory variables is extreme relative to other observations.

💡 In particular, for observation i ,

$$\text{leverage} = h_{ii} = \frac{1}{n} + \frac{(x_i - \bar{x})^2}{\sum (x_j - \bar{x})^2}$$

💡 Properties:

$$\frac{1}{n} \leq h_{ii} \leq 1$$

$$\sum h_{ii} = \text{tr}(H) = \text{rank}(H) = p + 1$$

💡 Also note

$$\hat{\mu} = H y \Rightarrow \hat{\mu}_i = h_{ii} y_i + \sum_{j \neq i} h_{ij} y_j$$

and so we can view the leverage as the weight of y_i 's contribution to the fitted value $\hat{\mu}_i$.

IDENTIFYING HIGH LEVERAGE CASES

💡 Residual plots are not useful in revealing high leverage points.

- as $h_{ii} \rightarrow 1$, $\text{Var}(e_i) \rightarrow 0$

- so high leverage observations' leverage is close to 0.

💡 Instead, we can just plot the leverage values.

- in R , we can use "hatvalues".

💡 We say observation i has high leverage if

$$h_{ii} > 2\bar{h} = \frac{2(p+1)}{n}.$$

INFLUENTIAL OBSERVATIONS

💡 We say an observation is "influential" if its removal from the line fit changes the fitted line (ie parameter estimates) considerably.

IDENTIFYING INFLUENTIAL OBSERVATIONS

💡 We can quantify influence using "Cook's distance":

$$D_i = \frac{(\hat{\mu} - \mu_{(i)})^T (\hat{\mu} - \mu_{(i)})}{\hat{\sigma}^2 (p+1)} = \frac{h_{ii}}{1-h_{ii}} \cdot \frac{d_i^2}{p+1}$$

where $\hat{\sigma}^2$ is the estimate of the variance from the model fit with the i^{th} observation included.

💡 We say an observation is strongly influential if

$$D_i \geq 1.$$

MODEL SELECTION

💡 Recall that

$$\hat{\sigma} = \sqrt{\frac{SS(\text{Res})}{n-(p+1)}}$$

💡 If we remove variables from the model, both $SS(\text{Res})$ & the df increase.

💡 If the increase in $SS(\text{Res})$ is small relative to the degrees of freedom gained, then $\hat{\sigma}$ will decrease, resulting in a more precise model.

💡 This is often the case for variables with large associated p-values.

ITERATIVE MODEL SELECTION

💡 Idea: build a reduced model by adding/removing variables one at a time, & refitting at each iteration until no more variables can be added/removed.

BACKWARD ELIMINATION

💡 Idea:

- 1 Fit all p variables;
- 2 Remove variable with largest p-value greater than some threshold (eg $\alpha=0.1$)
- 3 Refit with remaining variables;
- 4 Repeat 1-2 until no more variables can be removed.

FORWARD SELECTION

💡 Idea:

- 1 Fit all p single variable SLR models;
- 2 Select variable with smallest p-value $< \alpha$;
- 3 Fit the $p-1$ 2-variable models that include the variable in 2;
- 4 Repeat 1-3, continuing to add variables at each iteration until no more variables can be added.

STEPWISE SELECTION

💡 Idea:

- 1 Begin with forward selection;
- 2 Employ both forward selection & backward elimination at each step until no more variables can be added/removed.

ADJUSTED R-SQUARED

💡 Motivation: Recall

$$R^2 = 1 - \frac{SS(\text{Res})}{SS(\text{Tot})}$$

Note R^2 will always increase when more variables are added regardless of whether the variables account for a significant amount of the variation of the response.

💡 Hence, we cannot use R^2 to compare model subsets with differing amounts of parameters.

💡 Thus, we can use the "adjusted R -squared":

$$R_{\text{adj}}^2 = 1 - \frac{SS(\text{Res})/(n-(p+1))}{SS(\text{Tot})/(n-1)}$$

💡 R_{adj}^2 will only increase if the variation accounted for by the added variables increases proportionally more than the degrees of freedom decrease through the model parameters.

💡 Note we can also write

$$R_{\text{adj}}^2 = 1 - \frac{\hat{\sigma}^2}{SS(\text{Tot})/(n-1)}$$

So a large R_{adj}^2 implies a small $\hat{\sigma}$ and v.v.

MALLOWS' C_p

For a k -variable model ($k=1, \dots, p$), we define

$$C_p = \frac{SS(Res)_k}{MS(Res)_p} + 2(k+1) - n$$

Note smaller C_p values relative to the number of variables are associated with more suitable models.

A k -variable model is preferred over the full model if

$$C_p \leq k+1.$$

MODEL SELECTION IN R

To select the best model in R, we can use the 'leaps' function, which selects the best model from all possible subsets based on the R^2_{adj} & C_p^2 criteria.

eg

```
> leaps(house[, -9], value, method=c('adjr'), nbest=2, names=names(house[-9]))
size stories baths rooms age lotsize basement garage
1 TRUE FALSE FALSE FALSE FALSE FALSE FALSE FALSE
1 FALSE FALSE TRUE FALSE FALSE FALSE FALSE FALSE
2 TRUE FALSE FALSE FALSE TRUE FALSE FALSE FALSE
2 TRUE TRUE FALSE FALSE FALSE FALSE FALSE FALSE
3 TRUE FALSE FALSE FALSE TRUE TRUE FALSE FALSE
3 TRUE FALSE FALSE FALSE TRUE FALSE FALSE TRUE
4 TRUE FALSE FALSE FALSE TRUE TRUE FALSE TRUE
4 TRUE TRUE FALSE FALSE TRUE TRUE FALSE FALSE
5 TRUE TRUE FALSE FALSE TRUE TRUE FALSE TRUE
5 TRUE FALSE FALSE TRUE TRUE TRUE FALSE TRUE
6 TRUE TRUE TRUE FALSE TRUE TRUE FALSE TRUE
6 TRUE TRUE FALSE TRUE TRUE TRUE FALSE TRUE
7 TRUE TRUE TRUE TRUE TRUE TRUE FALSE TRUE
7 TRUE TRUE TRUE TRUE TRUE TRUE TRUE TRUE
8 TRUE TRUE TRUE TRUE TRUE TRUE TRUE TRUE

$adjr2
[1] 0.549407 0.317135 0.572833 0.560069 0.618481 0.593374 0.644686 0.6251132
[9] 0.651625 0.641082 0.648770 0.648308 0.645221 0.645063 0.641430
```

```
> leaps(house[, -9], value, method=c('Cp'), nbest=1, names=names(house[-9]))
size stories baths rooms age lotsize basement garage
1 TRUE FALSE FALSE FALSE FALSE FALSE FALSE FALSE
2 TRUE FALSE FALSE FALSE TRUE FALSE FALSE FALSE
3 TRUE FALSE FALSE FALSE TRUE TRUE FALSE FALSE
4 TRUE FALSE FALSE FALSE TRUE TRUE FALSE TRUE
5 TRUE TRUE FALSE FALSE TRUE TRUE FALSE TRUE
6 TRUE TRUE TRUE FALSE TRUE TRUE FALSE TRUE
7 TRUE TRUE TRUE TRUE TRUE TRUE FALSE TRUE
8 TRUE TRUE TRUE TRUE TRUE TRUE TRUE TRUE

$Cp
[1] 27.150656 21.556807 10.144146 4.137270 3.327332 5.096306 7.027277 9.000000
```

Note

- ① the full model has a higher R^2 value, as it has more parameters; but
- ② the reduced model has a higher R^2_{adj} value & a lower $\hat{\sigma}$ value.

INTERACTION

We say there is "interaction" between x_j, x_k if the effect of x_j on the response depends on the value of x_k .

To account for a possible interaction effect, we can include the term $x_j x_k$ in our model.

FITTING LINEAR MODELS TO TIME SERIES DATA

Q1: "Time series data" is data $\{y_t\}$, where y_t denotes the value of the response at time t ($t=1, 2, \dots$).

Q2: Variation in time series is caused by:

- ① "Seasonal" - repeat regularly over fixed intervals
- ② "Trend" - persistent increase/decrease in μ_t as t increases
- ③ "Cyclical" - oscillations repeating over irregular intervals
- ④ "Irregular/random" - other sources of variation unaccounted by ①-③.

AUTOCORRELATION FUNCTION / ACF

Q1: The "auto-correlation function" is

$$r_k = \frac{\sum_{t=k+1}^n (y_t - \bar{y})(y_{t-k} - \bar{y})}{\sum_{t=1}^n (y_t - \bar{y})^2}$$

and quantifies the degree of similarity between y_t & the "lagged" version of itself, y_{t-k} .

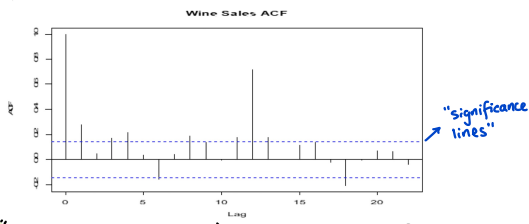
Q2: Note

- ① $|r_k| \leq 1 \quad \forall k;$ &
- ② r_k is unitless.

CORRELOGRAM

Q1: A "correlogram" is a plot of r_k over k (for $k=1, 2, \dots$).

les)



Q2: We can use the `acf` function in R to plot this.

SEASONAL COMPONENT

Q: Idea: We can add model parameters to account for the seasonal variation in the data, which might exacerbate any trends.

TREND COMPONENT

Q: Idea: We can account for the trend component in the model by adding appropriate time variables (eg t, t^2 , etc)

Chapter 3:

Logistic Regression

💡 Idea: Suppose we want to model a binary response variable

$$Y_i = \begin{cases} 1, & \text{if } i^{\text{th}} \text{ observation is successful} \\ 0, & \text{otherwise} \end{cases}$$

where $P(Y_i = 1) = \pi_i$.

💡 We can model $\mu_i = E(Y_i) = \pi_i$ using

$$\pi_i = \frac{e^{x_i^T \beta}}{1 + e^{x_i^T \beta}}$$

$$\Leftrightarrow \log\left(\frac{\pi_i}{1 - \pi_i}\right) = x_i^T \beta$$

where the LHS of the second equation is called the logit function.

💡 In R, we can fit a logistic regression model using

`> glm(y ~ x1 + ... + xp, family = binomial(link = 'logit'))`

INTERPRETATION OF LOGISTIC REGRESSION PARAMETER ESTIMATES

💡 By manipulating the equation

$$\log\left(\frac{\pi}{1 - \pi}\right) = x^T \beta = \beta_0 + \beta_1 x_1 + \dots + \beta_p x_p$$

we can get

$$e^{\beta_j} = \frac{\pi_{x_{j+1}} / (1 - \pi_{x_{j+1}})}{\pi / (1 - \pi)}$$

where $\pi_{x_{j+1}}$ is the value of $E(Y_i)$

after increasing one unit of x_j .

💡 Thus, we can interpret this "odds ratio" as the multiplicative change in the odds of success for a change in one unit of x_j , after accounting for the other variables.

DISTRIBUTION OF $\hat{\beta}_j$

💡 Let $\hat{\beta}_j$ be the MLE of β_j .

Then, for sufficiently large n , we can show that

$$\left(\frac{\hat{\beta}_j - \beta_j}{SE(\hat{\beta}_j)}\right)^2 \sim \chi_1^2$$

$$\Leftrightarrow \frac{\hat{\beta}_j - \beta_j}{SE(\hat{\beta}_j)} \sim N(0, 1)$$

$H_0: \beta_j = 0$ TEST

💡 To test $H_0: \beta_j = 0$, we can use the "Wald test statistic"

$$z = \frac{\hat{\beta}_j}{SE(\hat{\beta}_j)}$$

which is $\sim N(0, 1)$ under H_0 .

💡 Note the test statistic & p-values associated with this are given in the 'summary' output of the glm model in R.

💡 Hence, a $100(1 - \alpha)\%$ CI for β_j is

$$\hat{\beta}_j \pm z_{1 - \frac{\alpha}{2}} SE(\hat{\beta}_j)$$

and a $100(1 - \alpha)\%$ CI for the odds ratio is

$$e^{\hat{\beta}_j} \pm z_{1 - \frac{\alpha}{2}} SE(\hat{\beta}_j)$$

DEVIANCE / D

⚡₁ The likelihood function for a logistic regression model with a binary response is

$$L(\beta|y) = \prod_{i=1}^n f(y_i; \pi_i) = \prod_{i=1}^n \pi_i^{y_i} (1-\pi_i)^{1-y_i}$$

with corresponding log-likelihood

$$\ell(\beta|y) = \sum_{i=1}^n (y_i \log \pi_i + (1-y_i) \log(1-\pi_i))$$

⚡₂ The "deviance" of the model is

$$D = 2(\ell(\text{saturated}) - \ell(\text{fitted}))$$

where

- ① $\ell(\text{fitted})$ is the log-likelihood of the fitted model (with MLE $\hat{\pi}_i$); &
- ② $\ell(\text{saturated})$ is the log-likelihood of the saturated model, the hypothetical model for which $p_i = n$ (with MLE $\hat{\pi}_i = y_i$).

⚡₃ In particular, the deviance for the binary response logistic regression model is

$$D = -2 \sum (y_i \log(\frac{\hat{\pi}_i}{1-\hat{\pi}_i}) + \log(1-y_i))$$

MODEL COMPARISON - DEVIANCE

⚡₁ To test whether either $x_{u_1}, \dots, x_{u_m}$ are associated with a positive response of y after accounting for the other variates, we can use the test

$$H_0: \beta_{u_1} = \dots = \beta_{u_m} = 0.$$

⚡₂ Our full model is

$$\beta_0 + \sum_{j=1}^p \beta_j x_j$$

and our reduced model is

$$\beta_0 + \sum_{j \notin \{u_1, \dots, u_m\}} \beta_j x_j$$

⚡₃ If we assume H_0 is true (ie the reduced model gives a better fit), then we can approximate

$$\begin{aligned} D_{\text{red}} - D_{\text{full}} &= 2 \log \left(\frac{L(\text{fitted})_{\text{full}}}{L(\text{fitted})_{\text{red}}} \right) \\ &\sim \chi^2_{df_{\text{red}} - df_{\text{full}}} \end{aligned}$$

MODEL COMPARISON - AKIAKE INFORMATION CRITERION / AIC

⚡₁ The "AIC" is defined as

$$AIC = 2(p+1) - 2\ell(\text{fitted})$$

⚡₂ The smaller the AIC, the better the fit of the model.