

Basic Statistical Analysis in R

Get Working Directory

- `getwd()`

Set Working Directory

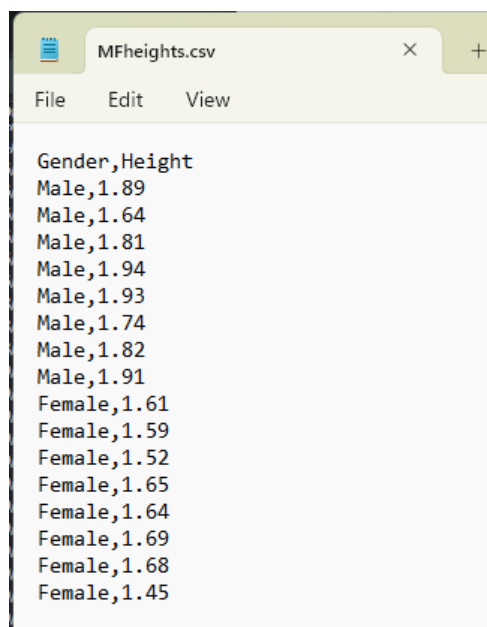
- `setwd("C:/Users/andyg/Desktop/R")` – for example

1. Descriptive Statistics

Height of Students

	A	B	C
1	Gender	Height	
2	Male	1.89	
3	Male	1.64	
4	Male	1.81	
5	Male	1.94	
6	Male	1.93	
7	Male	1.74	
8	Male	1.82	
9	Male	1.91	
10	Female	1.61	
11	Female	1.59	
12	Female	1.52	
13	Female	1.65	
14	Female	1.64	
15	Female	1.69	
16	Female	1.68	
17	Female	1.45	
18			

Save as “MFheights.csv”



Import into R

- `MFheights <- read.csv("MFheights.csv", header = TRUE, colClasses = c("factor", "numeric"))`

Look at data

- `summary(MFheights)`
- `summary(MFheights$Height)`

```
> summary(MFheights)
 1..Gender      Height
Female:8   Min.    :1.590
Male  :8   1st Qu.:1.688
          Median :1.775
          Mean   :1.776
          3rd Qu.:1.883
          Max.   :1.940
> summary(MFheights$Height)
   Min. 1st Qu.  Median    Mean 3rd Qu.    Max.
 1.590  1.688   1.775   1.776   1.883   1.940
```

Can also look individually.

- `mean(MFheights$Height)`
- `median(MFheights$Height)`
- `sd(MFheights$Height)`
- `var(MFheights$Height)`
- `range(MFheights$Height)`
- `quantile(MFheights$Height)`
- `IQR(MFheights$Height)`
- `min(MFheights$Height)`
- `max(MFheights$Height)`

- `table(MFheights$Height)`

- `boxplot(MFheights$Height)`
- `hist(MFheights$Height)`

2. Correlation

Relationship between ability in maths and sports.

	A	B	C	D	E	F
1	ID	Sport	Maths			
2	1	15	24			
3	2	19	29			
4	3	12	20			
5	4	14	27			
6	5	9	16			
7	6	15	28			
8	7	12	28			
9	8	14	23			
10	9	15	21			
11	10	19	29			
12						
13						
14						
15						

Save as “maths-sport.csv”

- `Maths.Sport <- read.csv("maths-sport.csv", header = TRUE, colClasses = c("numeric", "numeric", "numeric"))`
- `plot(Maths.Sport$Maths, Maths.Sport$Sport)`
- `cor(Maths.Sport$Maths, Maths.Sport$Sport)`
- `cor.test(Maths.Sport$Maths, Maths.Sport$Sport)`

```
> cor(Maths.Sport$Maths, Maths.Sport$Sport)
[1] 0.725103
> cor.test(Maths.Sport$Maths, Maths.Sport$Sport)

Pearson's product-moment correlation

data: Maths.Sport$Maths and Maths.Sport$Sport
t = 2.9782, df = 8, p-value = 0.01765
alternative hypothesis: true correlation is not equal to 0
95 percent confidence interval:
 0.1756842 0.9300985
sample estimates:
cor
0.725103
```

Yes, there is a strong significant relationship ($r = 0.725$, $p < 0.05$).

3. Regression

Can we predict reading ability from IQ score? What is the relationship?

	A	B	C	D	E	F	G
1	ID	IQ	Reading				
2	1	98	54				
3	2	115	74				
4	3	124	81				
5	4	97	59				
6	5	135	88				
7	6	105	70				
8	7	117	78				
9	8	101	73				
10	9	110	84				
11	10	127	86				
12							
13							
14							

Save as "IQ-Reading.csv"

- `IQ.Reading <- read.csv("IQ-Reading.csv", header = TRUE, colClasses = c("numeric", "numeric", "numeric"))`
- `cor.test(IQ.Reading$Reading, IQ.Reading$IQ)`
- `lm(IQ.Reading$Reading ~ IQ.Reading$IQ)`

```
> cor.test(IQ.Reading$Reading, IQ.Reading$IQ)

Pearson's product-moment correlation

data:  IQ.Reading$Reading and IQ.Reading$IQ
t = 4.626, df = 8, p-value = 0.001697
alternative hypothesis: true correlation is not equal to 0
95 percent confidence interval:
 0.4829915 0.9646213
sample estimates:
cor
0.8531664

> lm(IQ.Reading$Reading ~ IQ.Reading$IQ)

Call:
lm(formula = IQ.Reading$Reading ~ IQ.Reading$IQ)

Coefficients:
(Intercept)  IQ.Reading$IQ
-8.8247      0.7398
```

Yes, there is a strong significant relationship ($r=0.853$, $p=0.0017$)

Therefore: $\text{Reading} = 0.74 \times \text{IQ} - 8.82$

Multiple Regression

Can we predict reading ability from IQ score + number of Books Available? What is the relationship?

	A	B	C	D	E	F
1	ID	IQ	Books	Reading		
2		1	98	5	54	
3		2	115	25	74	
4		3	124	37	81	
5		4	97	10	59	
6		5	135	90	88	
7		6	105	80	70	
8		7	117	88	78	
9		8	101	73	73	
10		9	110	80	84	
11		10	127	100	86	
12						
13						

Save as "IQ-Books-Reading.csv"

- `IQ.Books.Reading <- read.csv("IQ-Books-Reading.csv", header = TRUE, colClasses = c("numeric", "numeric", "numeric", "numeric"))`
- `model <- lm(IQ.Books.Reading$Reading ~ IQ.Books.Reading$IQ + IQ.Books.Reading$Books)`
- `summary(model)`

```
R Console

> model <- lm(IQ.Books.Reading$Reading ~ IQ.Books.Reading$IQ + IQ.Books.Reading$Books)
> summary(model)

Call:
lm(formula = IQ.Books.Reading$Reading ~ IQ.Books.Reading$IQ +
    IQ.Books.Reading$Books)

Residuals:
    Min       1Q   Median       3Q      Max
-5.230 -2.939 -1.188  2.856  7.862

Coefficients:
            Estimate Std. Error t value Pr(>|t|)
(Intercept)   6.47837    14.58373   0.444  0.6703
IQ.Books.Reading$IQ  0.53111     0.14093   3.769  0.0070 **
IQ.Books.Reading$Books 0.14046     0.05118   2.745  0.0287 *
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 4.624 on 7 degrees of freedom
Multiple R-squared:  0.8689,    Adjusted R-squared:  0.8315
F-statistic: 23.2 on 2 and 7 DF,    p-value: 0.0008151
```

Yes, there is a very strong relationship ($R^2=0.869$, $p < 0.01$)

Therefore: $\text{Reading} = 6.478 + .531 \times \text{IQ} + .140 \times \text{Books Available}$

We can add more variables. For example: Male or Female (1=Male, 2=Female).

	A	B	C	D	E	F
1	ID	IQ	Books	Sex	Reading	
2	1	98	5	1	54	
3	2	115	25	2	74	
4	3	124	37	2	81	
5	4	97	10	1	59	
6	5	135	90	1	88	
7	6	105	80	1	70	
8	7	117	88	2	78	
9	8	101	73	2	73	
10	9	110	80	1	84	
11	10	127	100	2	86	
12						
13						

Save as "IQ-Books-Sex-Reading.csv"

- `IQ.Books.Sex.Reading <- read.csv("IQ-Books-Sex-Reading.csv", header = TRUE, colClasses = c("numeric", "numeric", "numeric", "numeric", "numeric"))`
- `model <- lm(IQ.Books.Sex.Reading$Reading ~ IQ.Books.Sex.Reading$IQ + IQ.Books.Sex.Reading$Books + IQ.Books.Sex.Reading$Sex)`
- `summary(model)`

```

R Console

> model <- lm(IQ.Books.Sex.Reading$Reading ~ IQ.Books.Sex.Reading$IQ + IQ.Books.Sex.Reading$Sex + IQ.Books.Sex.Reading$Books)
> summary(model)

Call:
lm(formula = IQ.Books.Sex.Reading$Reading ~ IQ.Books.Sex.Reading$IQ +
    IQ.Books.Sex.Reading$Sex + IQ.Books.Sex.Reading$Books)

Residuals:
    Min       1Q   Median       3Q      Max
-4.6713 -2.7087 -0.6339  1.9184  8.6998

Coefficients:
            Estimate Std. Error t value Pr(>|t|)
(Intercept)    6.3863    15.3608   0.416  0.6921
IQ.Books.Sex.Reading$IQ    0.5079     0.1542   3.294  0.0165 *
IQ.Books.Sex.Reading$Sex    1.8091     3.2472   0.557  0.5976
IQ.Books.Sex.Reading$Books  0.1405     0.0539   2.606  0.0403 *
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 4.87 on 6 degrees of freedom
Multiple R-squared:  0.8754,    Adjusted R-squared:  0.8131
F-statistic: 14.05 on 3 and 6 DF,  p-value: 0.004031

> |

```

Again there is a strong significant correlation between the proposed model of reading ability score

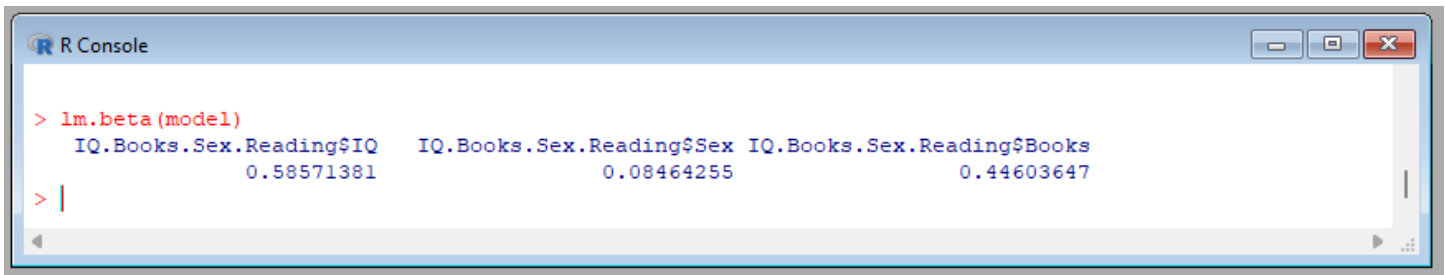
$$\text{Reading} = 6.39 + .5 \times \text{IQ} + .14 \times \text{Books} + 1.8 \times \text{Sex}$$

Note that Sex is a categorical variable with two levels (Males = 1 & Females = 2).

In this case the regression weighting (1.8) is added to the predicted Reading value. So for females 1.8 (1.8 x 1) would be added and for males 3.6 (1.8 x 2) would be added.

Which variables are important?

Install and load the “QuantPsych” package, and use the function “lm.beta” on your model.

An R Console window titled "R Console" with standard window controls (minimize, maximize, close). The console shows the execution of the `lm.beta(model)` function. The output is a table with three columns: `IQ.Books.Sex.Reading$IQ`, `IQ.Books.Sex.Reading$Sex`, and `IQ.Books.Sex.Reading$Books`. The corresponding values are 0.58571381, 0.08464255, and 0.44603647. A prompt `> |` is visible on the line below the output.

```
> lm.beta(model)
  IQ.Books.Sex.Reading$IQ  IQ.Books.Sex.Reading$Sex  IQ.Books.Sex.Reading$Books
0.58571381           0.08464255           0.44603647
> |
```

You now have the standardised regression coefficients, which you can now compare.

IQ: 0.58571381

Sex: 0.08464255

Books Available: 0.44603647

And you can see that the most important predictor is IQ, followed closely by Books Available. Sex is not contributing very much.

4. t-Test

Is there a difference between oral announcement of objects and memory?

	A	B	C	D	E	F	G	H
1	ID	Method	Recall					
2	1	Announcement	5					
3	2	Announcement	7					
4	3	Announcement	8					
5	4	Announcement	9					
6	5	Announcement	5					
7	6	Announcement	4					
8	7	Announcement	9					
9	8	Announcement	8					
10	9	Announcement	7					
11	10	Announcement	4					
12	11	No Announcement	3					
13	12	No Announcement	8					
14	13	No Announcement	5					
15	14	No Announcement	6					
16	15	No Announcement	3					
17	16	No Announcement	5					
18	17	No Announcement	9					
19	18	No Announcement	3					
20	19	No Announcement	3					
21	20	No Announcement	4					
22	21	No Announcement	7					
23	22	No Announcement	4					
24								

Save as “Announcement.csv”

- `Announcement <- read.csv("Announcement.csv", header = TRUE, colClasses = c("numeric", "factor", "numeric"))`
- `summary(Announcement)`
- `boxplot(Announcement$Recall ~ Announcement$Method)`

`t.test(numeric variable ~ dichotamous variable)`

OR `t.test(numeric variable1, numeric variable2)`

- `t.test(Announcement$Recall ~ Announcement$Method)`

```
> t.test(Announcement$Recall~Announcement$Method)

Welch Two Sample t-test

data:  Announcement$Recall by Announcement$Method
t = 1.8526, df = 19.69, p-value = 0.07898
alternative hypothesis: true difference in means between group Announcement and $
95 percent confidence interval:
 -0.2033249  3.4033249
sample estimates:
mean in group Announcement mean in group No Announcement
              6.6              5.0
```

No, there is no significant difference between oral announcement of objects and memory ($p = 0.08$).

5. ANOVA – One-Way

Is there a difference between amount of alcohol consumed and stopping distance?

	A	B	C	D	E	F	G
1	ID	alcohol	stopping				
2	1	No Alcohol	31.2				
3	2	No Alcohol	27.4				
4	3	No Alcohol	29.8				
5	4	No Alcohol	33.5				
6	5	No Alcohol	34				
7	6	One Unit	31.2				
8	7	One Unit	33.7				
9	8	One Unit	34.3				
10	9	One Unit	39.2				
11	10	One Unit	36				
12	11	Two Units	42.1				
13	12	Two Units	40.6				
14	13	Two Units	43.4				
15	14	Two Units	37.9				
16	15	Two Units	39				
17	16	Three Units	46.7				
18	17	Three Units	45.1				
19	18	Three Units	43.2				
20	19	Three Units	41.7				
21	20	Three Units	40.3				
22							

Stopping distance in metres.

Save as "Stopping.csv"

- `boxplot(Stopping$stopping ~ Stopping$alcohol)`

`aov(dependent numerical variable ~ factor)`

- `anova1 <- aov(Stopping$stopping ~ Stopping$alcohol)`
- `summary(anova1)`

```
> boxplot(Stopping$stopping~Stopping$alcohol)
> anova <- aov(Stopping$stopping ~ Stopping$alcohol)
> summary(anova)
```

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
Stopping\$alcohol	3	456.1	152.04	21.92	6.52e-06 ***
Residuals	16	111.0	6.94		

```
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
> |
```

Yes ($p < 0.05$). Therefore there is a significant difference in stopping distance according to amount of alcohol drunk.

Where are the differences?

- `sum <- TukeyHSD(anova)`
- `sum`

```
> TukeyHSD(anova)
Tukey multiple comparisons of means
 95% family-wise confidence level

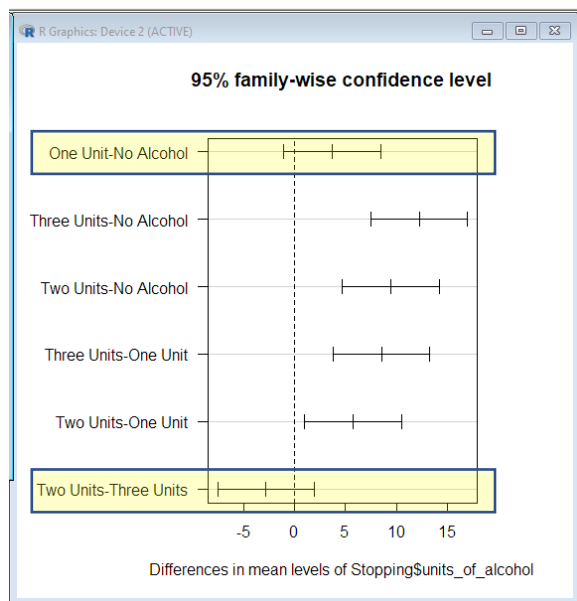
Fit: aov(formula = Stopping$stopping ~ Stopping$alcohol)

$`Stopping$alcohol`
              diff        lwr        upr      p adj
One Unit-No Alcohol    3.70 -1.0654663  8.465466 0.1596059
Three Units-No Alcohol 12.22  7.4545337 16.985466 0.0000091
Two Units-No Alcohol   9.42  4.6545337 14.185466 0.0001901
Three Units-One Unit    8.52  3.7545337 13.285466 0.0005430
Two Units-One Unit      5.72  0.9545337 10.485466 0.0161514
Two Units-Three Units  -2.80 -7.5654663  1.965466 0.3649771
```

- `plot(sum, las=1)`

Change margins

- `par(mar=c(5,10,5,5))`

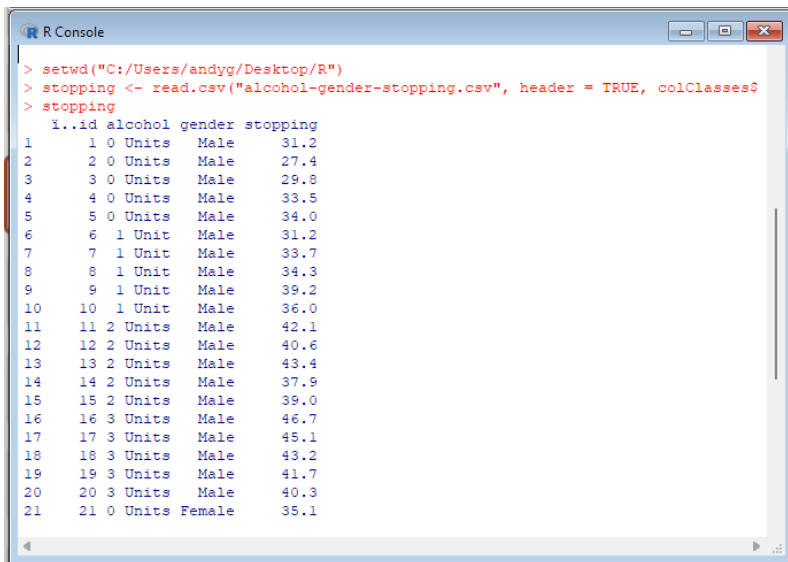


No significant difference when lines cross zero.

6. ANOVA – Two-Ways

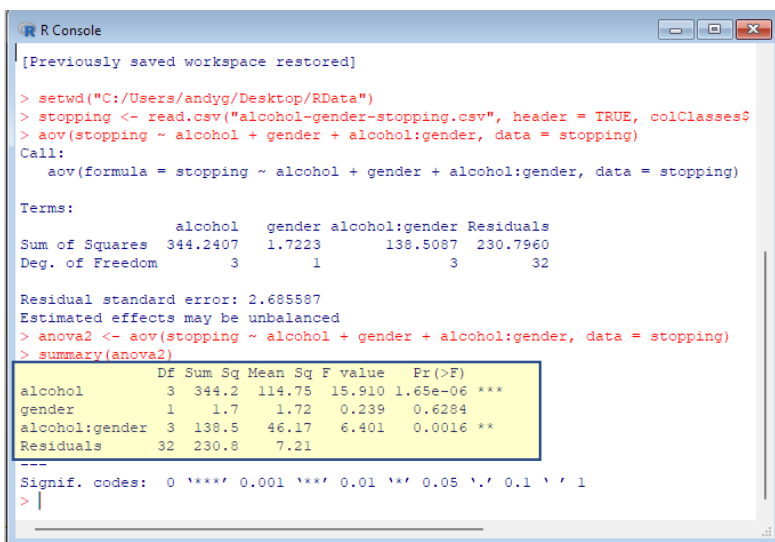
Is the effect of amount of alcohol consumed on stopping distance different for males and females?

- `stopping <- read.csv("alcohol-gender-stopping.csv", header = TRUE, colClasses = c("numeric", "factor", "factor", "numeric"))`



```
> setwd("C:/Users/andyg/Desktop/R")
> stopping <- read.csv("alcohol-gender-stopping.csv", header = TRUE, colClasses = c("numeric", "factor", "factor", "numeric"))
> stopping
  1..id alcohol gender stopping
1      1 0 Units   Male    31.2
2      2 0 Units   Male    27.4
3      3 0 Units   Male    29.8
4      4 0 Units   Male    33.5
5      5 0 Units   Male    34.0
6      6 1 Unit    Male    31.2
7      7 1 Unit    Male    33.7
8      8 1 Unit    Male    34.3
9      9 1 Unit    Male    39.2
10     10 1 Unit    Male    36.0
11     11 2 Units   Male    42.1
12     12 2 Units   Male    40.6
13     13 2 Units   Male    43.4
14     14 2 Units   Male    37.9
15     15 2 Units   Male    39.0
16     16 3 Units   Male    46.7
17     17 3 Units   Male    45.1
18     18 3 Units   Male    43.2
19     19 3 Units   Male    41.7
20     20 3 Units   Male    40.3
21     21 0 Units   Female   35.1
```

- `anova2 = aov(stopping ~ alcohol + gender + alcohol:gender, data = stopping)`
- `summary(anova2)`



```
[Previously saved workspace restored]
> setwd("C:/Users/andyg/Desktop/RData")
> stopping <- read.csv("alcohol-gender-stopping.csv", header = TRUE, colClasses = c("numeric", "factor", "factor", "numeric"))
> aov(stopping ~ alcohol + gender + alcohol:gender, data = stopping)
Call:
aov(formula = stopping ~ alcohol + gender + alcohol:gender, data = stopping)

Terms:
      alcohol      gender alcohol:gender Residuals
Sum of Squares 344.2407    1.7223      138.5087   230.7960
Deg. of Freedom      3         1          3        32

Residual standard error: 2.685587
Estimated effects may be unbalanced
> anova2 <- aov(stopping ~ alcohol + gender + alcohol:gender, data = stopping)
> summary(anova2)
```

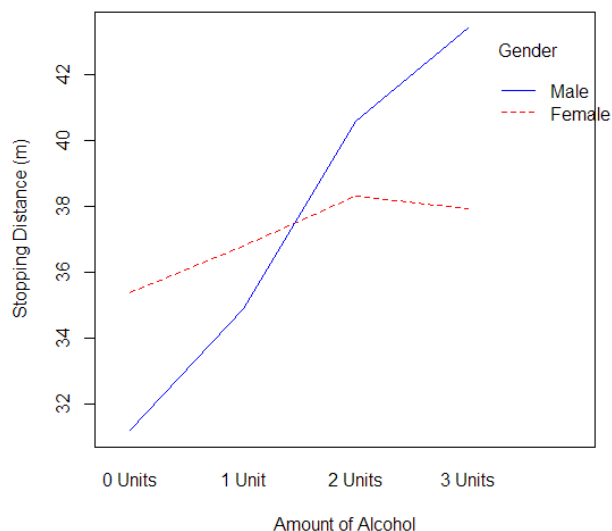
	Df	Sum Sq	Mean Sq	F value	Pr(>F)
alcohol	3	344.2	114.75	15.910	1.65e-06 ***
gender	1	1.7	1.72	0.239	0.6284
alcohol:gender	3	138.5	46.17	6.401	0.0016 **
Residuals	32	230.8	7.21		

```
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
>
```

alcohol:gender - Interaction effect = whether the influence of alcohol on stopping distance depends on whether you are male or female – significant ($p=0.0016$)

- `interaction.plot(stopping$alcohol, stopping$gender, stopping$stopping, col=c("red", "blue"), trace.label="Gender", xlab="Amount of Alcohol", ylab="Stopping Distance (m)", main="Stopping Distance According to Alcohol Consumption")`

Stopping Distance According to Alcohol Consumption



- `model_aov <- aov(anova1)`
- `TukeyHSD(model_aov)`

```
R Console
Male-Female 0.415 -1.31488 2.14488 0.628413
$`alcohol:gender`
      diff      lwr      upr      p adj
1 Unit:Female-0 Units:Female 1.42 -4.08199387 6.92199387 0.9895070
2 Units:Female-0 Units:Female 2.96 -2.54199387 8.46199387 0.6606975
3 Units:Female-0 Units:Female 2.58 -2.92199387 8.08199387 0.7917676
0 Units:Male-0 Units:Female -4.18 -9.68199387 1.32199387 0.2481314
1 Unit:Male-0 Units:Female -0.48 -5.98199387 5.02199387 0.9999911
2 Units:Male-0 Units:Female 5.24 -0.26199387 10.74199387 0.0710472
3 Units:Male-0 Units:Female 8.04 2.53800613 13.54199387 0.0010074
2 Units:Female-1 Unit:Female 1.54 -3.96199387 7.04199387 0.9832375
3 Units:Female-1 Unit:Female 1.16 -4.34199387 6.66199387 0.9969137
0 Units:Male-1 Unit:Female -5.60 -11.10199387 -0.09800613 0.0436885
1 Unit:Male-1 Unit:Female -1.90 -7.40199387 3.60199387 0.9478426
2 Units:Male-1 Unit:Female 3.82 -1.68199387 9.32199387 0.3513162
3 Units:Male-1 Unit:Female 6.62 1.11800613 12.12199387 0.0097756
3 Units:Female-2 Units:Female -0.38 -5.88199387 5.12199387 0.9999982
0 Units:Male-2 Units:Female -7.14 -12.64199387 -1.63800613 0.0043336
1 Unit:Male-2 Units:Female -3.44 -8.94199387 2.06199387 0.4817462
2 Units:Male-2 Units:Female 2.28 -3.22199387 7.78199387 0.8754904
3 Units:Male-2 Units:Female 5.08 -0.42199387 10.58199387 0.0874134
0 Units:Male-3 Units:Female -6.76 -12.26199387 -1.25800613 0.0078740
1 Unit:Male-3 Units:Female -3.06 -8.56199387 2.44199387 0.6235640
2 Units:Male-3 Units:Female 2.66 -2.84199387 8.16199387 0.7660547
```

7. ANCOVA

Does age affect the stopping distance according to the amount of alcohol drunk?

```
R Console
>
> stopping <- read.csv("alcohol-age-stopping.csv", header = TRUE, colClasses = $
> stopping
  i..id alcohol age stopping
1      1 0 Units 25.0    31.2
2      2 0 Units 24.0    27.4
3      3 0 Units 23.5    29.8
4      4 0 Units 24.7    33.5
5      5 0 Units 24.2    34.0
6      6 1 Unit 24.5    31.2
7      7 1 Unit 26.0    33.7
8      8 1 Unit 27.0    34.3
9      9 1 Unit 22.0    39.2
10     10 1 Unit 25.0    36.0
11     11 2 Units 24.3    42.1
12     12 2 Units 24.6    40.6
13     13 2 Units 28.0    43.4
14     14 2 Units 27.0    37.9
15     15 2 Units 26.0    39.0
16     16 3 Units 23.0    46.7
17     17 3 Units 25.0    45.1
18     18 3 Units 23.0    43.2
19     19 3 Units 28.0    41.7
20     20 3 Units 31.0    40.3
>
```

- `ancova<-aov(stopping$stopping ~ stopping$alcohol+stopping$age)`
- `summary(ancova)`

```
R Console
> ancova<-aov(stopping$stopping~stopping$alcohol+stopping$age)
>
> summary(ancova)
              Df Sum Sq Mean Sq F value    Pr(>F)
stopping$alcohol  3  456.1   152.04    25.61 3.76e-06 ***
stopping$age      1   21.9    21.91     3.69  0.074 .
Residuals       15   89.1     5.94
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
>
```

- `TukeyHSD(ancova,"stopping$alcohol")`

```
R Console
>
> TukeyHSD(ancova,"stopping$alcohol")
Tukey multiple comparisons of means
 95% family-wise confidence level

Fit: aov(formula = stopping$stopping ~ stopping$alcohol + stopping$age)

$`stopping$alcohol`
      diff      lwr      upr    p adj
1 Unit-0 Units  3.70 -0.7417801  8.14178 0.1196368
2 Units-0 Units  9.42  4.9782199 13.86178 0.0001053
3 Units-0 Units 12.22  7.7782199 16.66178 0.0000052
2 Units-1 Unit   5.72  1.2782199 10.16178 0.0100402
3 Units-1 Unit   8.52  4.0782199 12.96178 0.0003028
```

- `library('emmeans')`
- `emmeans(ancova, 'alcohol')`

```
R Console

      diff      lwr      upr      p adj
1 Unit-0 Units  3.70 -0.7417801  8.14178 0.1196368
2 Units-0 Units  9.42  4.9782199 13.86178 0.0001053
3 Units-0 Units 12.22  7.7782199 16.66178 0.0000052
2 Units-1 Unit   5.72  1.2782199 10.16178 0.0100402
3 Units-1 Unit   8.52  4.0782199 12.96178 0.0003028
3 Units-2 Units  2.80 -1.6417801  7.24178 0.3039307

Warning message:
In replications(paste("~", xx), data = mf) :
  non-factors ignored: stopping$age
> emmeans(ancova, 'alcohol')
Error in emmeans(ancova, "alcohol") : could not find function "emmeans"
> library('emmeans')
> emmeans(ancova, 'alcohol')
  alcohol emmean   SE df lower.CL upper.CL
0 Units   30.6 1.13 15    28.2    33.0
1 Unit    34.7 1.10 15    32.3    37.0
2 Units   41.0 1.11 15    38.6    43.3
3 Units   43.8 1.11 15    41.4    46.2

Confidence level used: 0.95
> |
```

After adjusting for age, a significant difference ($F(3, 15) = 25.61, p < .001$) and was found between the amount of alcohol drunk and the stopping distance. Post hoc comparisons using the Tukey test were carried out. No significant difference was found between stopping distances with No Alcohol and One Unit of alcohol ($p = 0.12$), with only a 3.42 difference in mean stopping distances, but the other differences were significant, with the mean difference in stopping distance ranging from 5.7 to 12.2 meters

An ANCOVA is a form of analysis that tells you whether your groups differ on a dependent variable when you have partialled out (removed the effects of) another variable called the covariate that has a relationship with the dependent variable. ANCOVA therefore adjusts the means on the covariate so that the mean covariate score is the same for all groups and therefore any changes in the DV are solely due to the treatment variable.

It is mostly used in pre-test-post-test designs where researchers want to partial out (remove or hold constant) the effect of a confounding variable. Using simple difference scores between pre and post-test does not achieve this. ANCOVA then answers the question, "What would the means of each group be on the DV if the means of the different groups had the same mean on the pre-test score?", i.e. it adjusts the DV means so that we can investigate the mean differences for post-test score only.

8. ANCOVA – Two-Way

A sample of male and female adults was selected and randomly allocated to four teaching groups using different teaching methods, after which they were given a test.

Group 1: Method 1

Group 2: Method 2

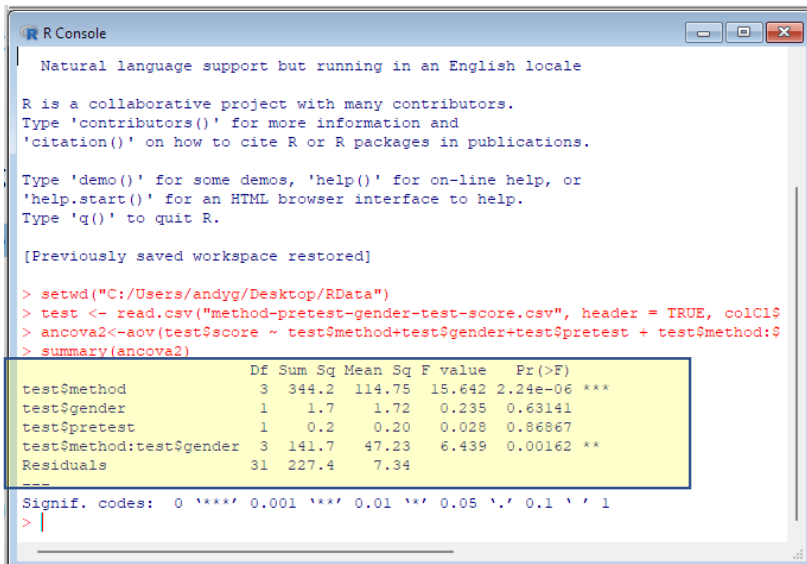
Group 3: Method 3

Group 4: Method 4

Two categorical variables – Method & Gender

Two continuous variables – Pre-Test Score & Final Test Score

- `test <- read.csv("method-pretest-gender-test-score.csv", header = TRUE, colClasses = c("factor", "factor", "numeric", "numeric"))`
- `ancova2 <- aov(test$score ~ test$method + test$gender + test$pretest + test$method:test$gender)`
- `summary(ancova2)`



```
R Console

Natural language support but running in an English locale

R is a collaborative project with many contributors.
Type 'contributors()' for more information and
'citation()' on how to cite R or R packages in publications.

Type 'demo()' for some demos, 'help()' for on-line help, or
'help.start()' for an HTML browser interface to help.
Type 'q()' to quit R.

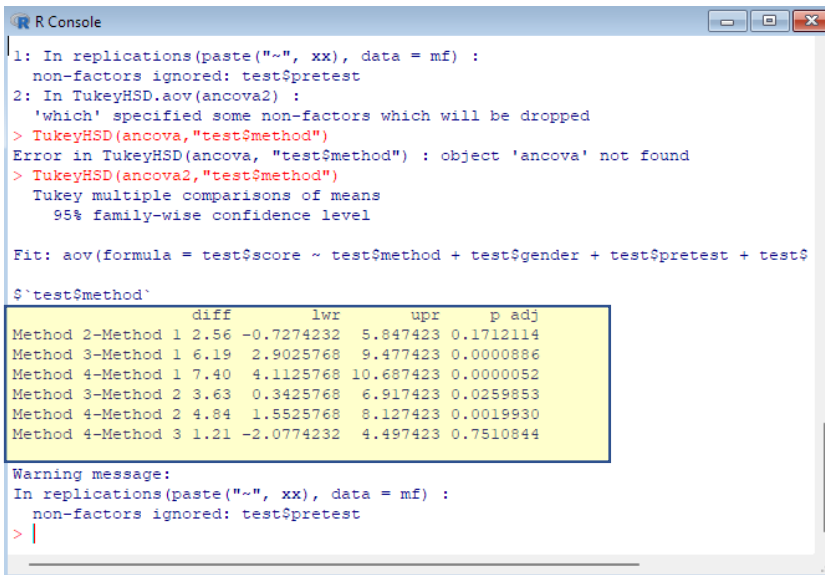
[Previously saved workspace restored]

> setwd("C:/Users/andyg/Desktop/RData")
> test <- read.csv("method-pretest-gender-test-score.csv", header = TRUE, colClasses = c("factor", "factor", "numeric", "numeric"))
> ancova2 <- aov(test$score ~ test$method + test$gender + test$pretest + test$method:test$gender)
> summary(ancova2)
```

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
test\$method	3	344.2	114.75	15.642	2.24e-06 ***
test\$gender	1	1.7	1.72	0.235	0.63141
test\$pretest	1	0.2	0.20	0.028	0.86867
test\$method:test\$gender	3	141.7	47.23	6.439	0.00162 **
Residuals	31	227.4	7.34		

```
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
> |
```

- `TukeyHSD(ancova2, "test$method")`



```
R Console

1: In replications(paste("~", xx), data = mf) :
  non-factors ignored: test$pretest
2: In TukeyHSD.aov(ancova2) :
  'which' specified some non-factors which will be dropped
> TukeyHSD(ancova, "test$method")
Error in TukeyHSD(ancova, "test$method") : object 'ancova' not found
> TukeyHSD(ancova2, "test$method")
  Tukey multiple comparisons of means
    95% family-wise confidence level

Fit: aov(formula = test$score ~ test$method + test$gender + test$pretest + test$method:test$gender)

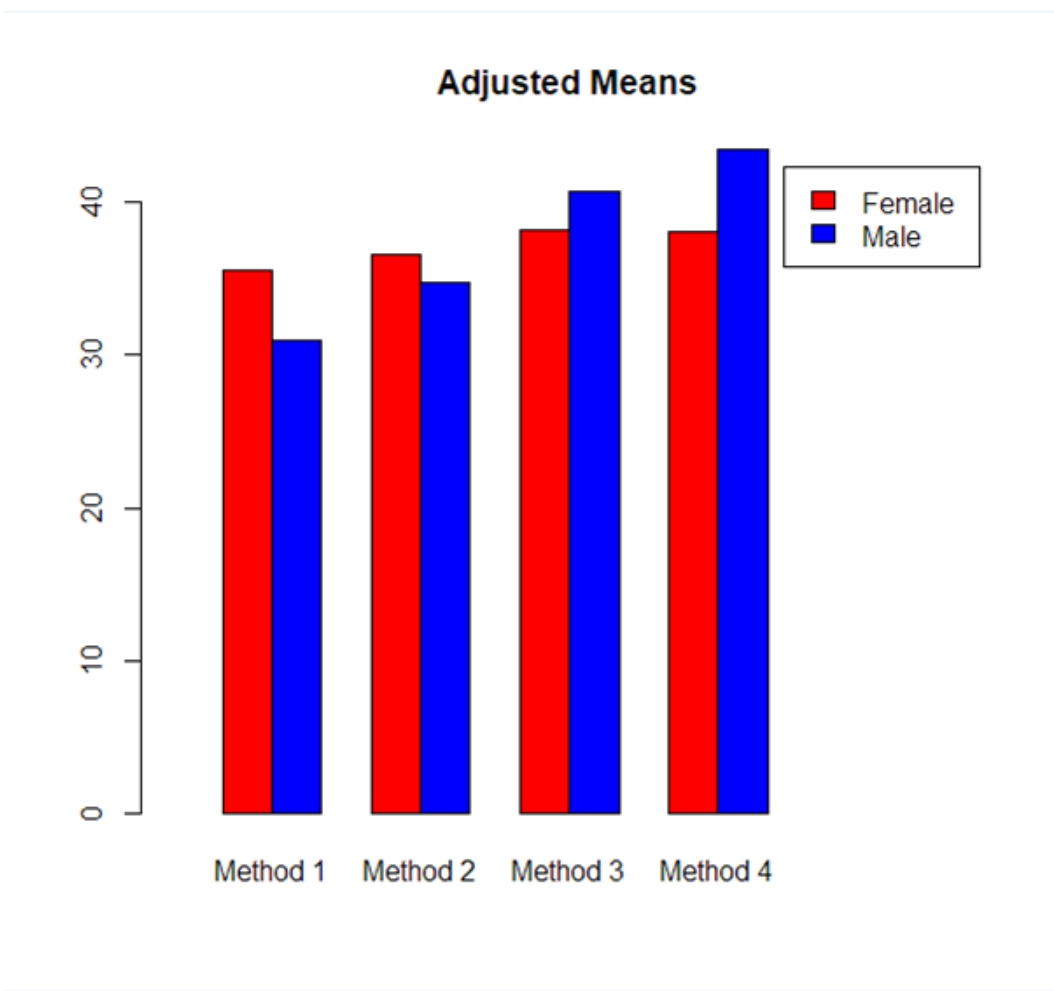
$`test$method`
      diff      lwr      upr      p adj
Method 2-Method 1 2.56 -0.7274232  5.847423 0.1712114
Method 3-Method 1 6.19  2.9025768  9.477423 0.0000886
Method 4-Method 1 7.40  4.1125768 10.687423 0.0000052
Method 3-Method 2 3.63  0.3425768  6.917423 0.0259853
Method 4-Method 2 4.84  1.5525768  8.127423 0.0019930
Method 4-Method 3 1.21 -2.0774232  4.497423 0.7510844

Warning message:
In replications(paste("~", xx), data = mf) :
  non-factors ignored: test$pretest
> |
```

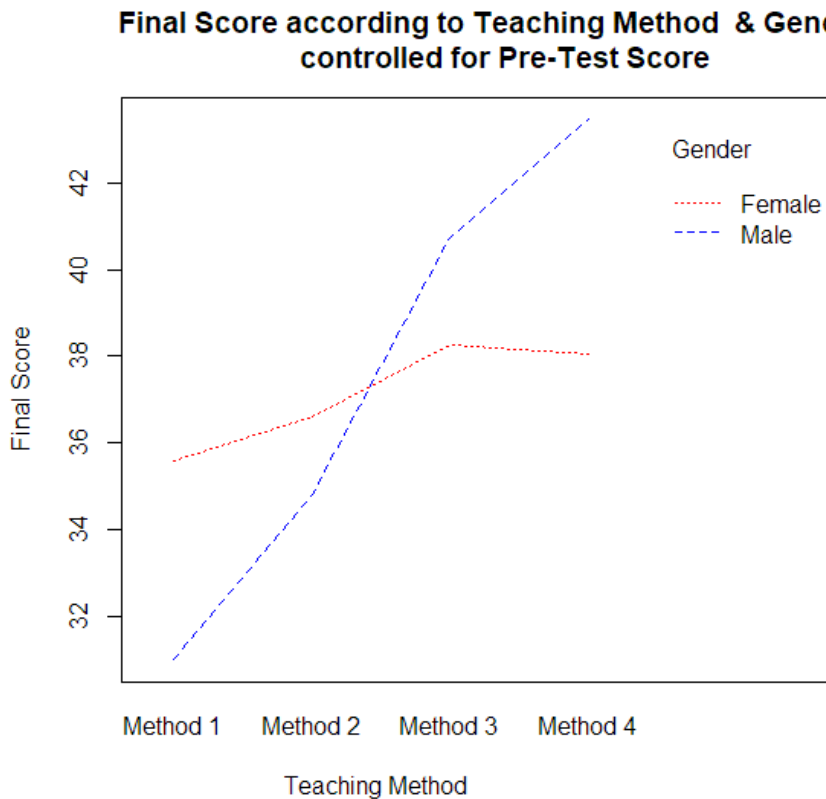
Using the “effects” package.

- library("effects")
- result <- aov(score ~ pretest + method*gender, data=test)
- effect_method = effect("method*gender",result)
- data.frame(effect_method)

```
R Console
>
> tapply(test$score, list(test$gender, test$method), mean)
      Method 1 Method 2 Method 3 Method 4
Female    35.36    36.78    38.32    37.94
Male      31.18    34.88    40.60    43.40
> library('effects')
> result <- aov(score ~ pretest + method*gender, data=test)
> effect_method = effect("method*gender",result)
> data.frame(effect_method)
  method gender      fit      se   lower   upper
1 Method 1 Female 35.59523 1.259696 33.02606 38.16440
2 Method 2 Female 36.59387 1.241804 34.06119 39.12655
3 Method 3 Female 38.24457 1.216322 35.76386 40.72527
4 Method 4 Female 38.05739 1.223496 35.56206 40.55273
5 Method 1 Male 30.99387 1.241804 28.46119 33.52655
6 Method 2 Male 34.80457 1.216322 32.32386 37.28527
7 Method 3 Male 40.71739 1.223496 38.22206 43.21273
8 Method 4 Male 43.45312 1.213768 40.97762 45.92861
> |
```



- `interaction.plot(means$method, means$gender, means$fit, col=c("red", "blue"), trace.label="Gender", xlab="Teaching Method", ylab="Final Score", main= "Final Score according to Teaching Method & Gender, \n controlled for Pre-Test Score")`



A 2 by 2 between-groups analysis of covariance was conducted to assess the effectiveness of four methods of teaching English for male and female participants. The independent variables were the teaching method (method 1, 2, 3 or 4) and gender. The dependent variable was scores on the final test, administered following completion of the intervention programs. Scores on the pre-test administered prior to the commencement of the programs were used as a covariate to control for individual differences.

After adjusting for pre-test scores, there was a significant interaction effect: $F(1, 31) = 6.439, p < .001$. These results suggest that males and females respond differently to the four methods. Males showed a more substantial increase on the final test score after participation the class using methods 1 and 2 than methods 3 & 4, whereas females did not differ so much.

9. Chi-Squared (χ^2) (Test for independence)

Is there a significant difference between the proportion of males and females who smoke?

	Smoke	Not Smoke
Male	18	24
Female	35	14

Input data

- `smokers <- matrix(c(18,35,24,14), ncol=2)`
- `rownames(smokers) <- c("Male", "Female")`
- `colnames(smokers) <- c("Smoke", "Not Smoke")`

- smokers

```
> smokers <- matrix(c(18,35,24,14), ncol=2)
> rownames(smokers) <- c("Male", "Female")
> colnames(smokers) <- c("Smoke", "Not Smoke")
> smokers
      Smoke Not Smoke
Male      18       24
Female    35       14
```

- `chisq.test(smokers)`

```
> chisq.test(smokers)
```

Pearson's Chi-squared test with Yates' continuity correction

data: smokers

X-squared = 6.4615, df = 1, p-value = 0.01102

OR

	A	B	C	D	E	F	G	H	I	J
24	24	Male	Not Smoke							
26	25	Male	Not Smoke							
27	26	Male	Not Smoke							
28	27	Male	Not Smoke							
29	28	Male	Not Smoke							
30	29	Male	Not Smoke							
31	30	Male	Not Smoke							
32	31	Male	Not Smoke							
33	32	Male	Not Smoke							
34	33	Male	Not Smoke							
35	34	Male	Not Smoke							
36	35	Male	Not Smoke							
37	36	Male	Not Smoke							
38	37	Male	Not Smoke							
39	38	Male	Not Smoke							
40	39	Male	Not Smoke							
41	40	Male	Not Smoke							
42	41	Male	Not Smoke							
43	42	Male	Not Smoke							
44	43	Female	Smoke							
45	44	Female	Smoke							
46	45	Female	Smoke							
47	46	Female	Smoke							
48	47	Female	Smoke							
49	48	Female	Smoke							
50	49	Female	Smoke							
51	50	Female	Smoke							
52	51	Female	Smoke							
53	52	Female	Smoke							

Save as "smoke.csv"

Import data into R

- `smoke <- read.csv("smoke.csv", header = TRUE, colClasses = c("numeric", "factor", "factor"))`
- `MFsmoke <- table(smoke$Sex, smoke$Smoke)`
- `chisq.test(MFsmoke)`

```
> MFsmoke <- table(smoke$Sex, smoke$Smoke)
> MFsmoke

      Not Smoke Smoke
Female      14    35
Male       24    18
> chisq.test(MFsmoke)

      Pearson's Chi-squared test with Yates' continuity correction

data:  MFsmoke
X-squared = 6.4615, df = 1, p-value = 0.01102
```

$p < 0.05$ - we can conclude that there is a significant difference between males and females in their smoking activity.