

Psychology 350: An introduction to R for Psychological Research

Week 5: Student's t and Fisher's F

William Revelle
Northwestern University
Evanston, Illinois USA

<https://personality-project.org/courses/350>



NORTHWESTERN
UNIVERSITY

Outline

The t-test

t and the central limit theorem

ANOVA

ANOVA: and Im
ANOVA

Linear Regression
Regression from the raw data

t.test demonstration with Student's data using cushny dataset

William Gossett, publishing under the name [Student \(1908\)](#) reported a small sample approximation (t) to the large sample z test. His first example was a data set on the different effect of optical isomers of hyoscyamine hydrobromide reported by [Cushny and Peebles \(1905\)](#). The sleep of 10 patients was measured without any drug and then following administration of D. and L isomers. The data from Cushny are available as the cushny data set.

Variable	Cntrl	drug1	drg2L	drg2R	delt1	dlt2L	dlt2R
1	0.60	1.3	2.50	2.10	0.70	1.90	1.50
2	3.00	1.4	3.80	4.40	-1.60	0.80	1.40
3	4.70	4.5	5.80	4.70	-0.20	1.10	0.00
4	5.50	4.3	5.60	4.80	-1.20	0.10	-0.70
5	6.20	6.1	6.10	6.70	-0.10	-0.10	0.50
6	3.20	6.6	7.60	8.30	3.40	4.40	5.10
7	2.50	6.2	8.00	8.20	3.70	5.50	5.70
8	2.80	3.6	4.40	4.30	0.80	1.60	1.50
9	1.10	1.1	5.70	5.80	0.00	4.60	4.70
10	2.90	4.9	6.30	6.40	2.00	3.40	3.50
Mean	3.25	4.0	5.58	5.57	0.75	2.33	2.32
Sd	1.78	2.1	1.66	1.91	1.79	2.00	2.27

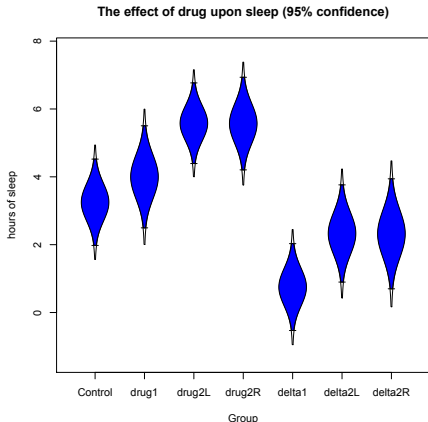
R code

```
error.bars(cushny, xlab="Group", ylab="hours of sleep",
           main="The effect of drug upon sleep (95% confidence)")
```

The cushny data set with error bars (Cushny and Peebles, 1905)

R code

```
error.bars(cushny,xlab="Group",ylab="hours of sleep",  
           main="The effect of drug upon sleep (95% confidence)")
```



We can show these data graphically using the `error.bars` function. We pass labels to the `x` and `y` axis using the `xlab` and `ylab` parameters, and then supply an appropriate figure title.

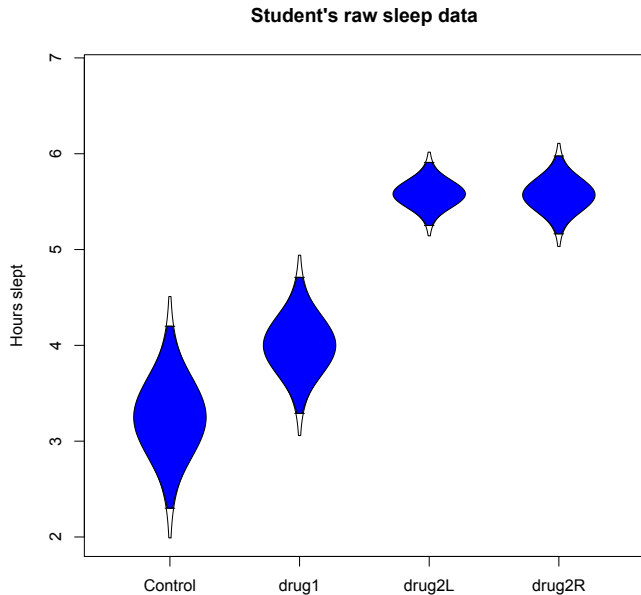
We will use these data to show how to do t-tests as well as the generalization to Analysis of Variance.

Student's t.test: As done by Student

R code

```
with(cushny,t.test(delta1)) #control versus drug 1 difference scores
with(cushny,t.test(delta2L)) #control versus drug2L difference scores
with(cushny,t.test(delta1,delta2L,paired=TRUE)) #difference of differences
```

```
> with(cushny,t.test(delta1)) #control versus drug 1 difference scores
  One Sample t-test
data:  delta1
t = 1.3257, df = 9, p-value = 0.2176
alternative hypothesis: true mean is not equal to 0
95 percent confidence interval:
 -0.5297804  2.0297804
sample estimates:
mean of x
 0.75
with(cushny,t.test(delta2L)) #control versus drug2L difference scores
  One Sample t-test
data:  delta2L
t = 3.6799, df = 9, p-value = 0.005076
alternative hypothesis: true mean is not equal to 0
95 percent confidence interval:
 0.8976775 3.7623225
sample estimates:
mean of x
 2.33
> with(cushny,t.test(delta1,delta2L,paired=TRUE)) #difference of differences
  Paired t-test
data:  delta1 and delta2L
t = -4.0621, df = 9, p-value = 0.002833
alternative hypothesis: true difference in means is not equal to 0
95 percent confidence interval:
 -2.4598858 -0.7001142
sample estimates:
```



Two ways of organizing the data: Wide versus long

We can take the wide format of the `cushny` data set and make it long.

```
cushny[c("delta1", "delta2L")]
```

	delta1	delta2L
1	0.7	1.9
2	-1.6	0.8
3	-0.2	1.1
4	-1.2	0.1
5	-0.1	-0.1
6	3.4	4.4
7	3.7	5.5
8	0.8	1.6
9	0.0	4.6
10	2.0	3.4

R code

```
long.sleep <-  
  stack(cushny[c("delta1", "delta2L")])  
long.sleep
```

	values	ind
1	0.7	delta1
2	-1.6	delta1
3	-0.2	delta1
4	-1.2	delta1
5	-0.1	delta1
6	3.4	delta1
7	3.7	delta1
8	0.8	delta1
9	0.0	delta1
10	2.0	delta1
11	1.9	delta2L
12	0.8	delta2L
13	1.1	delta2L
14	0.1	delta2L
15	-0.1	delta2L
16	4.4	delta2L
17	5.5	delta2L
18	1.6	delta2L
19	4.6	delta2L
20	3.4	delta2L

R code

```
long.sleep <-
+   stack(cushny[c("delta1",
+                   "delta2L")])
```

```
> long.sleep
  values    ind
1    0.7 delta1
2   -1.6 delta1
3   -0.2 delta1
4   -1.2 delta1
5   -0.1 delta1
6    3.4 delta1
7    3.7 delta1
8    0.8 delta1
9    0.0 delta1
10   2.0 delta1
11   1.9 delta2L
12   0.8 delta2L
13   1.1 delta2L
14   0.1 delta2L
15  -0.1 delta2L
16   4.4 delta2L
17   5.5 delta2L
18   1.6 delta2L
19   4.6 delta2L
20   3.4 delta2L
```

R code

```
t.test(values ~ ind, data=long.sleep)
```

```
Welch Two Sample t-test
data: values by ind
t = -1.8608, df = 17.776, p-value = 0.07939
alternative hypothesis: true difference in means is not e
95 percent confidence interval:
 -3.3654832  0.2054832
sample estimates:
mean in group delta1 mean in group delta2L
          0.75              2.33
```

But, the data were paired

R code

```
t.test(values ~ ind, data=long.sleep,
       paired=TRUE)
```

```
data: values by ind
t = -4.0621, df = 9, p-value = 0.002833
alternative hypothesis: true difference in means is not e
95 percent confidence interval:
 -2.4598858 -0.7001142
sample estimates:
mean of the differences
          -1.58
```


t.test demonstration with Student's data (from the sleep dataset)

Sleep data set is
just 2 columns of
cushny

R code

sleep

```
> sleep
  extra group ID
1    0.7      1  1
2   -1.6      1  2
3   -0.2      1  3
4   -1.2      1  4
5   -0.1      1  5
6    3.4      1  6
7    3.7      1  7
8    0.8      1  8
9    0.0      1  9
10   2.0      1 10
11   1.9      2  1
12   0.8      2  2
13   1.1      2  3
14   0.1      2  4
15  -0.1      2  5
16   4.4      2  6
17   5.5      2  7
18   1.6      2  8
19   4.6      2  9
20   3.4      2 10
```

R code

```
with(sleep,t.test(extra~group))
with(sleep,t.test(extra~group,var.equal=TRUE))
```

Welch Two Sample t-test

```
data: extra by group
t = -1.8608, df = 17.776, p-value = 0.07939 <-- default value
t = -1.8608, df = 18, p-value = 0.07919. <- equal variances
alternative hypothesis: true difference in means is not equal to 0
95 percent confidence interval:
 -3.3654832  0.2054832
sample estimates:
mean in group 1 mean in group 2
      0.75          2.33
```

But the data were actually paired. Do it for a paired t-test

R code

```
with(sleep,t.test(extra~group,paired=TRUE))
```

Paired t-test

```
data: extra by group
t = -4.0621, df = 9, p-value = 0.002833
alternative hypothesis: true difference in means is not equal to 0
95 percent confidence interval:
 -2.4598858 -0.7001142
sample estimates:
mean of the differences
      -1.58
```

What are these things called p values and confidence intervals?

1. Data are estimate of population values
2. Each sample statistic will differ from sample to sample
3. How much these sample statistics differ is of great concern
4. The central limit theorem suggests that the distributions of means will be lawful and have a *normal* distribution
5. But the range of our estimates depends upon the sample size.
6. We can compare observed distributions with *normal theory* but for smaller samples, these estimates are too liberal.

The Central Limit Theorem

1. One of the most amazing statistical results is known as the central limit theorem. For any distribution with finite variance, if we take progressively bigger samples from the population, the distribution of the means of those samples will tend towards the normal distribution. (The CLT does not apply to distributions with infinite variance.)
2. See the demonstration in [week 5 handout](#)

t and the estimate of variance

1. As shown by Gosset, small samples will underestimate the population standard deviation.
2. This is why he developed the t distribution.
3. We can show this by simulation.

The original Gossett/Student data set

1. The classic data set used by Gossett (publishing as [Student \(1908\)](#) for the introduction of the t-test.
2. The design was a within subjects study with hours of sleep in a control condition compared to those in 3 drug conditions.
 - Drug1 was 06mg of L Hscyamine,
 - Drug 2L and Drug2R were said to be .6 mg of Left and Right isomers of Hyoscine.
 - However as discussed by Zabell (2008) these were not optical isomers.
 - The deta1, delta2L and delta2R are changes from the baseline control.

The original Gossett/Student data set

R code

```
cushny
describe(cushy)
desribe(sleep) #from the sleep data set compares delta21 to delta2L
```

	Control	drug1	drug2L	drug2R	delta1	delta2L	delta2R
1	0.6	1.3	2.5	2.1	0.7	1.9	1.5
2	3.0	1.4	3.8	4.4	-1.6	0.8	1.4
3	4.7	4.5	5.8	4.7	-0.2	1.1	0.0
4	5.5	4.3	5.6	4.8	-1.2	0.1	-0.7
5	6.2	6.1	6.1	6.7	-0.1	-0.1	0.5
6	3.2	6.6	7.6	8.3	3.4	4.4	5.1
7	2.5	6.2	8.0	8.2	3.7	5.5	5.7
8	2.8	3.6	4.4	4.3	0.8	1.6	1.5
9	1.1	1.1	5.7	5.8	0.0	4.6	4.7
10	2.9	4.9	6.3	6.4	2.0	3.4	3.5

	vars	n	mean	sd	median	trimmed	mad	min	max	range	skew	kurtosis	se
Control	1	10	3.25	1.78	2.95	3.21	1.63	0.6	6.2	5.6	0.20	-1.23	0.56
drug1	2	10	4.00	2.10	4.40	4.04	2.59	1.1	6.6	5.5	-0.25	-1.67	0.66
drug2L	3	10	5.58	1.66	5.75	5.66	1.41	2.5	8.0	5.5	-0.30	-0.99	0.53
drug2R	4	10	5.57	1.91	5.30	5.66	1.56	2.1	8.3	6.2	-0.09	-1.07	0.60
delta1	5	10	0.75	1.79	0.35	0.67	1.56	-1.6	3.7	5.3	0.42	-1.30	0.57
delta2L	6	10	2.33	2.00	1.75	2.24	2.45	-0.1	5.5	5.6	0.28	-1.66	0.63
delta2R	7	10	2.32	2.27	1.50	2.28	2.59	-0.7	5.7	6.4	0.23	-1.68	0.72

	vars	n	mean	sd	median	trimmed	mad	min	max	range	skew	kurtosis	se
extra	1	20	1.54	2.02	0.95	1.47	1.56	-1.6	5.5	7.1	0.39	-1.09	0.45
group*	2	20	1.50	0.51	1.50	1.50	0.74	1.0	2.0	1.0	0.00	-2.10	0.11
ID*	3	20	5.50	2.95	5.50	5.50	3.71	1.0	10.0	9.0	0.00	-1.40	0.66

R code

```
with(cushny, t.test(drug1, drug2L, paired=TRUE)) #within subjects
```

Paired t-test

```
data: drug1 and drug2L
t = -4.0621, df = 9, p-value = 0.002833
alternative hypothesis: true mean difference is not equal to 0
95 percent confidence interval:
 -2.4598858 -0.7001142
sample estimates:
mean difference
      -1.58
```

Analysis of Variance as special case of linear model

1. aov provides a wrapper to lm for fitting linear models to balanced or unbalanced experimental designs.
2. The main difference from lm is in the way print, summary and so on handle the fit: this is expressed in the traditional language of the analysis of variance rather than that of linear models.
3. If the formula contains a single Error term, this is used to specify error strata, and appropriate models are fitted within each error stratum.
4. The formula can specify multiple responses.
5. aov is designed for balanced designs, and the results can be hard to interpret without balance: beware that missing values in the response(s) will likely lose the balance.
6. If there are two or more error strata, the methods used are statistically inefficient without balance, and it may be better to use lme in package nlme.

Several examples of AOV

Several examples are available.

1. One is found by `?aov` and is taken from [Venables and Ripley \(2002\)](#). This is an example of the effect of Nitrogen (N) Potassium (K) and Phosphate (P) on the growth of Peas. (Remember that the people who developed ANOVA were agronomists.).
2. Another example is the effect of gender and drug on alertness.
3. First, we will examine the `sleep` data set, using `aov` instead of the `t.test`.

aov of the sleep data set: compare with the t.test results

R code

> sleep

```
extra group ID
1      0.7      1  1
2     -1.6      1  2
3     -0.2      1  3
4     -1.2      1  4
5     -0.1      1  5
6      3.4      1  6
7      3.7      1  7
8      0.8      1  8
9      0.0      1  9
10     2.0      1 10
11     1.9      2  1
12     0.8      2  2
13     1.1      2  3
14     0.1      2  4
15    -0.1      2  5
16     4.4      2  6
17     5.5      2  7
18     1.6      2  8
19     4.6      2  9
20     3.4      2 10
```

R code

```
#independent subjects
summary(aov(extra ~ group, data=sleep))
```

```
> summary(aov(extra ~ group, data=sleep))
              Df Sum Sq Mean Sq F value Pr(>F)
group          1  12.48   12.482   3.463 0.0792 .
Residuals     18   64.89    3.605
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
t = -1.8608, df = 17.776, p-value = 0.07939. <--
t = -1.8608, df = 18, p-value = 0.07919. <- equal variances
```

R code

```
#correlated subjects
summary(aov(extra~group + Error(ID), data=sleep))
```

```
> summary(aov(extra~group + Error(ID), data=sleep))
```

Error: ID

```
              Df Sum Sq Mean Sq F value Pr(>F)
Residuals     9   58.08    6.453
```

Error: Within

```
              Df Sum Sq Mean Sq F value Pr(>F)
group          1  12.482   12.482   16.5 0.00283 **
Residuals     9    6.808    0.756
```

```
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
t = -4.0621, df = 9, p-value = 0.002833. <---
```

aov: an example of chemicals upon the growth of peas.

R code

npk #from Venables

	block	N	P	K	yield
1	1	0	1	1	49.5
2	1	1	1	0	62.8
3	1	0	0	0	46.8
4	1	1	0	1	57.0
5	2	1	0	0	59.8
6	2	1	1	1	58.5
7	2	0	0	1	55.5
8	2	0	1	0	56.0
9	3	0	1	0	62.8
10	3	1	1	1	55.8
11	3	1	0	0	69.5
12	3	0	0	1	55.0
13	4	1	0	0	62.0
14	4	1	1	1	48.8
15	4	0	0	1	45.5
16	4	0	1	0	44.2
17	5	1	1	0	52.0
18	5	0	0	0	51.5
19	5	1	0	1	49.8
20	5	0	1	1	48.8
21	6	1	0	1	57.2
22	6	1	1	0	59.0
23	6	0	1	1	53.2
24	6	0	0	0	56.0

Several models

R code

```
mod1 <- aov(yield ~ N,data=npk)
mod2 <- aov(yield ~ N+ P + N*P,data=npk)
mod2a <- aov(yield ~N*P,data=npk)
mod3 <- aov(yield ~ N*P*K,data=npk)
mod4 <- aov(yield ~ block + N*P*K,data=npk)
```

> summary(mod1)

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
N	1	189.3	189.28	6.061	0.0221 *
Residuals	22	687.1	31.23		

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

> summary(mod4)

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
block	5	343.3	68.66	4.447	0.01594 *
N	1	189.3	189.28	12.259	0.00437 **
P	1	8.4	8.40	0.544	0.47490
K	1	95.2	95.20	6.166	0.02880 *
N:P	1	21.3	21.28	1.378	0.26317
N:K	1	33.1	33.13	2.146	0.16865
P:K	1	0.5	0.48	0.031	0.86275
Residuals	12	185.3	15.44		

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

R code

```
mod1 <- aov(yield ~ N,data=npk)
mod2 <- aov(yield ~ N+ P + N*P,data=npk)
mod2a <- aov(yield ~N*P,data=npk)
mod3 <- aov(yield ~ N*P*K,data=npk)
mod4 <- aov(yield ~ block + N*P*K,data=npk)
```

```
summary(mod1)
```

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
N	1	189.3	189.28	6.061	0.0221 *
Residuals	22	687.1	31.23		

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

```
> summary(mod2)
```

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
N	1	189.3	189.28	5.758	0.0263 *
P	1	8.4	8.40	0.256	0.6187
N:P	1	21.3	21.28	0.647	0.4305
Residuals	20	657.4	32.87		

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

R code

```
mod2a <- aov(yield ~N*P,data=npk) #same as before
```

```
summary(mod3)
```

	Df	Sum Sq	Mean Sq	F	value	Pr(>F)	
N	1	189.3	189.28	6.161	0.0245	*	
P	1	8.4	8.40	0.273	0.6082		
K	1	95.2	95.20	3.099	0.0975	.	
N:P	1	21.3	21.28	0.693	0.4175		
N:K	1	33.1	33.14	1.078	0.3145		
P:K	1	0.5	0.48	0.016	0.9019		
N:P:K	1	37.0	37.00	1.204	0.2887		

ANOVA and the linear model

1. aov handles the data as a linear model but prints it out in the conventional anova table
2. We can see this if we compare an lm approach to the aov approach
3. The big difference is handling interactions
4. First compare aov without interactions to lm

aov and lm for main effects

R code

```
mod1 <- aov(yield ~ N + P + K, data = npk)
summary(mod1)
lm1 <- lm(yield ~ N + P + K, data = npk)
summary(lm1)
```

```
mod1 <- aov(yield ~ N + P + K, data = npk)
> summary(mod1)
```

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
N	1	189.3	189.28	6.488	0.0192 *
P	1	8.4	8.40	0.288	0.5974
K	1	95.2	95.20	3.263	0.0859 .
Residuals	20	583.5	29.17		

```
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
> lm1 <- lm(yield ~ N + P + K, data = npk)
> summary(lm1)
Call:
lm(formula = yield ~ N + P + K, data = npk)
Residuals:
    Min       1Q   Median       3Q      Max
-9.2667 -3.6542  0.7083  3.4792  9.3333
Coefficients:
            Estimate Std. Error t value Pr(>|t|)
(Intercept)   54.650      2.205   24.784  <2e-16 ***
N1             5.617      2.205    2.547   0.0192 *
P1            -1.183      2.205   -0.537   0.5974
K1            -3.983      2.205   -1.806   0.0859 .
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
```

aov and lm differ in the way they treat interactions

1. aov treats all data as 'factors' and has a design matrix of orthogonal contrasts
2. lm treats in data as numeric and takes product terms for interactions
3. zero centering the data make these two approaches equivalent (if there are only two levels)
4. To zero center, use the 'scale' function
5. But first have to make the data 'numeric'
6. One way to do this is the `char2numeric`
7. `scale` returns a matrix, we need to make it a data frame

Centering the npk data set

R code

```
npk.n <- char2numeric(npk, flag=FALSE)
npk.c <- scale(npk.n, scale=FALSE)
npc.c <- as.data.frame(npk.c)
summary(lm(yield ~ N*P * K, data=npk.c))
```

Call:

```
lm(formula = yield ~ N * P * K, data = npk.c)
```

Residuals:

Min	1Q	Median	3Q	Max
-10.133	-4.133	1.250	3.125	8.467

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	-1.137e-15	1.131e+00	0.000	1.0000
N	5.617e+00	2.263e+00	2.482	0.0245 *
P	-1.183e+00	2.263e+00	-0.523	0.6082
K	-3.983e+00	2.263e+00	-1.760	0.0975 .
N:P	-3.767e+00	4.526e+00	-0.832	0.4175
N:K	-4.700e+00	4.526e+00	-1.038	0.3145
P:K	5.667e-01	4.526e+00	0.125	0.9019
N:P:K	9.933e+00	9.052e+00	1.097	0.2887

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

Residual standard error: 5.543 on 16 degrees of freedom

Multiple R-squared: 0.4391, Adjusted R-squared: 0.1937

F-statistic: 1.789 on 7 and 16 DF, p-value: 0.1586

Compare this to the aov model

R code

```
mod3 <- aov(yield ~ N*P*K, data=npk)
```

```
summary(mod3)
```

	Df	Sum Sq	Mean Sq	F value	Pr(>F)	
N	1	189.3	189.28	6.161	0.0245	*
P	1	8.4	8.40	0.273	0.6082	
K	1	95.2	95.20	3.099	0.0975	.
N:P	1	21.3	21.28	0.693	0.4175	
N:K	1	33.1	33.14	1.078	0.3145	
P:K	1	0.5	0.48	0.016	0.9019	
N:P:K	1	37.0	37.00	1.204	0.2887	

What if we don't center?

R code

```
mod3a <- lm(yield ~ N * P * K, data=npk)
summary(mod3a)
```

Call:

```
lm(formula = yield ~ N * P * K, data = npk)
```

Residuals:

Min	1Q	Median	3Q	Max
-10.133	-4.133	1.250	3.125	8.467

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	51.4333	3.2002	16.072	2.7e-11 ***
N1	12.3333	4.5258	2.725	0.015 *
P1	2.9000	4.5258	0.641	0.531
K1	0.5667	4.5258	0.125	0.902
N1:P1	-8.7333	6.4004	-1.365	0.191
N1:K1	-9.6667	6.4004	-1.510	0.150
P1:K1	-4.4000	6.4004	-0.687	0.502
N1:P1:K1	9.9333	9.0515	1.097	0.289

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

Residual standard error: 5.543 on 16 degrees of freedom

Multiple R-squared: 0.4391, Adjusted R-squared: 0.1937

F-statistic: 1.789 on 7 and 16 DF, p-value: 0.1586

ANOVA

Analysis of Variance: Another example

aov is designed for balanced designs, and the results can be hard to interpret without balance: beware that missing values in the response(s) will likely lose the balance.

R code

```
datafilename="http://personality-project.org/r/datasets/R.appendix2.data"
data.ex2=read.file(datafilename) #read the data into a data.frame
data.ex2 #show the data
```

data.ex2

Observation	Gender	Dosage	Alertness
1	1	m	a
2	2	m	a
3	3	m	a
4	4	m	a
5	5	m	b
6	6	m	b
7	7	m	b
8	8	m	b
9	9	f	a
10	10	f	a
11	11	f	a
12	12	f	a
13	13	f	b
14	14	f	b
15	15	f	b
16	16	f	b

R code

```
#do the analysis of variance
aov.ex2 = aov(Alertness~Gender*Dosage,data=data.ex2)
summary(aov.ex2) #show the summary table
```

Call:

```
summary(aov.ex2) #show the summary table
```

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
Gender	1	76.56	76.56	2.952	0.111
Dosage	1	5.06	5.06	0.195	0.666
Gender:Dosage	1	0.06	0.06	0.002	0.962
Residuals	12	311.25	25.94		

ANOVA

Analysis of Variance

Do the analysis of variances and then show the table of results.

R code

```
#do the analysis of variance
aov.ex2 <- aov(Alertness ~ Gender * Dosage, data=data.ex2)

summary(aov.ex2)           #show the summary table
aov.ex2. #This shows the coefficients
```

```
>aov.ex2 <- aov(Alertness ~ Gender * Dosage, data=data.ex2)
> summary(aov.ex2)           #show the summary table
      Df Sum Sq Mean Sq F value Pr(>F)
Gender    1  76.56    76.56   2.952  0.111
Dosage     1   5.06     5.06   0.195  0.666
Gender:Dosage 1   0.06     0.06   0.002  0.962
Residuals 12 311.25    25.94
```

```
aov(formula = Alertness ~ Gender * Dosage, data = data.ex2)
```

Terms:

	Gender	Dosage	Gender:Dosage	Residuals
Sum of Squares	76.5625	5.0625	0.0625	311.2500
Deg. of Freedom	1	1	1	12

Residual standard error: 5.092887

Estimated effects may be unbalanced

Show the results table

R code

```
print(model.tables(aov.ex2, "means"), digits=3)
```

```
> print(model.tables(aov.ex2, "means"), digits=3)
```

```
Tables of means
```

```
Grand mean
```

```
14.0625
```

```
Gender
```

```
Gender
```

```
      f      m
```

```
16.25 11.88
```

```
Dosage
```

```
Dosage
```

```
      a      b
```

```
13.50 14.62
```

```
Gender:Dosage
```

```
      Dosage
```

```
Gender a      b
```

```
      f 15.75 16.75
```

```
      m 11.25 12.50
```

Analysis of Variance: Within subjects

1. Somewhat more complicated because we need to convert “wide” data.frames to “long” or “narrow” data.frames.
2. This can be done by using the stack function. Some data sets are already in the long format.
3. A detailed discussion of how to work with repeated measures designs is at <http://personality-project.org/r/r.anova.html> and at <http://personality-project.org/r>
4. See also the tutorial by Jason French at <http://jason-french.com/tutorials/repeatedmeasures.html>
5. Many within subject designs can be treated as multi-level designs. For a discussion of analyzing multilevel data (particularly for personality dynamics), see <http://personality-project.org/revelle/publications/rw.paid.17.final.pdf>

ANOVA

Analysis of variance within subjects: Getting and showing the data

R code

```
filename="http://personality-project.org/r/datasets/R.appendix5.data"
data.ex5=read.file(filename)      #read the data into a data.frame
headTail(data.ex5,6,12)          #show the data (first 6, last 12)
```

	Obs	Subject	Gender	Dosage	Task	Valence	Recall
1	1	A	M	A	F	Neg	8
2	2	A	M	A	F	Neu	9
3	3	A	M	A	F	Pos	5
4	4	A	M	A	C	N eg	7
5	5	A	M	A	C	Neu	9
6	6	A	M	A	C	Pos	10
...	...	<NA>	<NA>	<NA>	<NA>	<NA>	...
97	97	Q	F	C	F	Neg	18
98	98	Q	F	C	F	Neu	17
99	99	Q	F	C	F	Pos	18
100	100	Q	F	C	C	Neg	17
101	101	Q	F	C	C	Neu	19
102	102	Q	F	C	C	Pos	19
103	103	R	F	C	F	Neg	19
104	104	R	F	C	F	Neu	17
105	105	R	F	C	F	Pos	19
106	106	R	F	C	C	Neg	22
107	107	R	F	C	C	Neu	21
108	108	R	F	C	C	Pos	20

Describe the data

R code

`describe(data.ex5)`

	vars	n	mean	sd	median	trimmed	mad	min	max	range	skew	kurtosis	se
Obs	1	108	54.50	31.32	54.5	54.50	40.03	1	108	107	0.00	-1.23	3.01
Subject*	2	108	9.50	5.21	9.5	9.50	6.67	1	18	17	0.00	-1.24	0.50
Gender*	3	108	1.50	0.50	1.5	1.50	0.74	1	2	1	0.00	-2.02	0.05
Dosage*	4	108	2.00	0.82	2.0	2.00	1.48	1	3	2	0.00	-1.53	0.08
Task*	5	108	1.50	0.50	1.5	1.50	0.74	1	2	1	0.00	-2.02	0.05
Valence*	6	108	2.00	0.82	2.0	2.00	1.48	1	3	2	0.00	-1.53	0.08
Recall	7	108	15.63	5.07	15.0	15.74	4.45	4	25	21	-0.13	-0.64	0.49

The * signify that the entries are not numerical, but rather categorical or logical.

ANOVA

Analysis of variance within subjects

R code

```
filename="http://personality-project.org/r/datasets/R.appendix5.data"
data.ex5=read.table(filename,header=TRUE) #read the data into a table
#do the anova
aov.ex5 = aov(Recall~(Task*Valence*Gender*Dosage)+Error(Subject/(Task*Valence))+
  (Gender*Dosage),data.ex5)
#look at the output
summary(aov.ex5)
```

Error: Subject

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
Gender	1	542.3	542.3	5.685	0.0345 *
Dosage	2	694.9	347.5	3.643	0.0580 .
Gender:Dosage	2	70.8	35.4	0.371	0.6976
Residuals	12	1144.6	95.4		

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

Error: Subject:Task

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
Task	1	96.33	96.33	39.862	3.87e-05 ***
Task:Gender	1	1.33	1.33	0.552	0.472
Task:Dosage	2	8.17	4.08	1.690	0.226
Task:Gender:Dosage	2	3.17	1.58	0.655	0.537
Residuals	12	29.00	2.42		

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
+ lots more

ANOVA

Analysis of variance within subjects output (continued)

Error: Subject:Valence

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
Valence	2	14.69	7.343	2.998	0.0688 .
Valence:Gender	2	3.91	1.954	0.798	0.4619
Valence:Dosage	4	20.26	5.065	2.068	0.1166
Valence:Gender:Dosage	4	1.04	0.259	0.106	0.9793
Residuals	24	58.78	2.449		

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

Error: Subject:Task:Valence

	Df	Sum Sq	Mean Sq	F value	Pr(>F)
Task:Valence	2	5.39	2.6944	1.320	0.286
Task:Valence:Gender	2	2.17	1.0833	0.531	0.595
Task:Valence:Dosage	4	2.78	0.6944	0.340	0.848
Task:Valence:Gender:Dosage	4	2.67	0.6667	0.327	0.857
Residuals	24	49.00	2.0417		

Multiple regression

1. Use the `sat.act` data set from *psych*
2. Do the linear model
3. Summarize the results

```
mod1 <- lm(SATV ~ education + gender + SATQ, data=sat.act)
> summary(mod1, digits=2)
```

Call:

```
lm(formula = SATV ~ education + gender + SATQ, data = sat.act)
```

Residuals:

Min	1Q	Median	3Q	Max
-372.91	-49.08	2.30	53.68	251.93

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	180.87348	23.41019	7.726	3.96e-14 ***
education	1.24043	2.32361	0.534	0.59363
gender	20.69271	6.99651	2.958	0.00321 **
SATQ	0.64489	0.02891	22.309	< 2e-16 ***

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

Residual standard error: 86.24 on 683 degrees of freedom
(13 observations deleted due to missingness)

Multiple R-squared: 0.4231, Adjusted R-squared: 0.4205
F-statistic: 167 on 3 and 683 DF, p-value: < 2.2e-16

Regression from the raw data

Zero center the data before examining interactions

```
> zsat <- data.frame(scale(sat.act, scale=FALSE))
> mod2 <- lm(SATV ~ education * gender * SATQ, data=zsat)
> summary(mod2)

Call:
lm(formula = SATV ~ education * gender * SATQ, data = zsat)
```

Residuals:

Min	1Q	Median	3Q	Max
-372.53	-48.76	3.33	51.24	238.50

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)	
(Intercept)	0.773576	3.304938	0.234	0.81500	
education	2.517314	2.337889	1.077	0.28198	
gender	18.485906	6.964694	2.654	0.00814	**
SATQ	0.620527	0.028925	21.453	< 2e-16	***
education:gender	1.249926	4.759374	0.263	0.79292	
education:SATQ	-0.101444	0.020100	-5.047	5.77e-07	***
gender:SATQ	0.007339	0.060850	0.121	0.90404	
education:gender:SATQ	0.035822	0.041192	0.870	0.38481	

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

Compare model 1 and model 2

Test the difference between the two linear models

```
> anova(mod1,mod2)
```

Analysis of Variance Table

Model 1: SATV ~ education + gender + SATQ

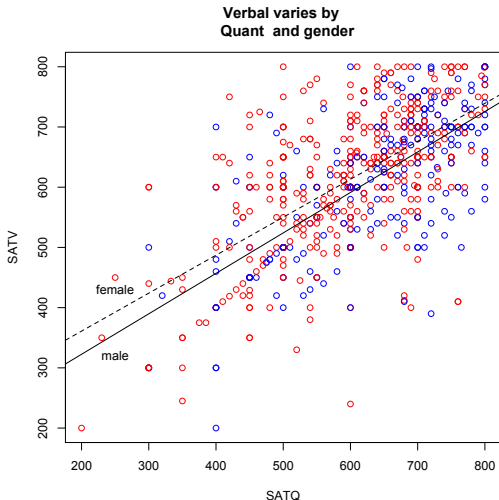
Model 2: SATV ~ education * gender * SATQ

	Res.Df	RSS	Df	Sum of Sq	F	Pr(>F)
1	683	5079984				
2	679	4870243	4	209742	7.3104	9.115e-06 ***

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.'

Regression from the raw data

Show the regression lines by gender



```
> with(sat.act, plot(SATV~SATQ,  
  col=c("blue", "red")[gender]))  
> by(sat.act, sat.act$gender,  
  function(x) abline  
    (lm(SATV~SATQ, data=x),  
    lty=c("solid", "dashed"))  
> title("Verbal varies by Quant  
  and gender")
```

Regression from the raw data

- Cushny, A. R. and Peebles, A. R. (1905). The action of optical isomer. ii. hyoscines. *The Journal of Physiology*, 32(501-510).
- Student (1908). The probable error of a mean. *Biometrika*, 6(1):1–25.
- Venables, W. N. and Ripley, B. D. (2002). *Modern Applied Statistics with S*. Springer, New York, fourth edition. ISBN 0-387-95457-0.