

17.8. How to report factor analysis ①

When reporting factor analysis we should provide our readers with form an informed opinion about what we've done. We should be for extracting factors and the method of rotation used. We should of the rotated factor loadings of all items and flag (in bold) values (I would personally choose .40, but see Section 17.4.6.2). We should percentage of variance that each factor explains and possibly the eigenvalues. Table 17.1 shows an example of such a table for the SAQ data (oblique rotation). We should also report the sample size in the title.

In my opinion, a table of factor loadings and a description of the factors is a minimum. You could consider (if it's not too large) including the table of communalities which someone could reproduce your analysis (should they want to). You should also mention sample size adequacy. For this example we might write something like:

- ✓ A principal axis factor analysis was conducted on the 23 item questionnaire (direct oblimin). The Kaiser–Meyer–Olkin measure verified the appropriateness for the analysis, $KMO = .93$ ('marvellous' according to Hatcher, 1999), and all KMO values for individual items were greater than the acceptable limit of .5 (Field, 2013). An initial analysis of the eigenvalues for each factor in the data. Four factors had eigenvalues greater than 1 and in combination explained 50.32% of the variance. The first factor was ambiguous and showed inflexions that would justify retaining it. We retained 4 factors because of the large sample size and the scree plot and Kaiser's criterion on this value. Table 17.1 shows the factor loadings after rotation. The items that cluster on the same factor represent a fear of statistics, factor 2 represents peer evaluation, factor 3 a fear of computers and factor 4 a fear of maths.

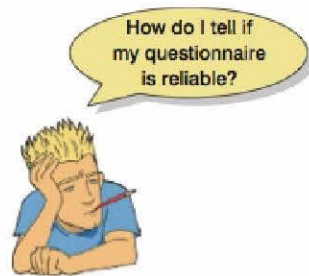
17.9. Reliability analysis ②

17.9.1. Measures of reliability ③

If you're using factor analysis to validate a questionnaire, it is useful to calculate the reliability of your scale.



SELF-TEST Thinking back to Chapter 1, what are the measures of reliability and test–retest reliability?



Reliability means that a measure (or in this case questionnaire) consistently reflect the construct that it is measuring. One that, other things being equal, a person should get the same score on a questionnaire if they complete it at two different points in time (this is called test-retest reliability). So, if a person who is familiar with SPSS and who scores highly on our SAQ should get a similar score if we tested them a month later (assuming they hadn't

CHAPTER 17 EXPLORATORY FACTOR ANALYSIS

TABLE 17.1 Summary of exploratory factor analysis results for the SPSS anxiety questionnaire (N = 2571)

Item	Rotated Factor Loadings		
	Fear of Statistics	Peer Evaluation	Fear of Computers
I wake up under my duvet thinking that I am trapped under a normal distribution	.54	-.04	.17
I can't sleep for thoughts of eigenvectors	.47	-.14	-.08
I weep openly at the mention of central tendency	.45	-.05	.17
I dream that Pearson is attacking me with correlation coefficients	.44	.08	.18
Standard deviations excite me	-.44	.32	-.05
Statistics makes me cry	.43	.10	.11
People try to tell you that SPSS makes statistics easier to understand but it doesn't	.41	-.04	.36
I don't understand statistics	.36	.05	.20
My friends are better at statistics than me	-.09	.56	-.02
My friends are better at SPSS than I am	.07	.47	-.11
My friends will think I'm stupid for not being able to cope with SPSS	-.18	.45	.04
If I'm good at statistics my friends will think I'm a nerd	.10	.35	.00
Everybody looks at me when I use SPSS	-.22	.34	-.08
I have little experience of computers	-.22	-.01	.86
SPSS always crashes when I try to use it	.18	-.01	.64
All computers hate me	.19	-.02	.56
I worry that I will cause irreparable damage because of my incompetence with computers	.08	-.04	.56
Computers have minds of their own and deliberately go wrong whenever I use them	.24	-.02	.47
Computers are useful only for playing games	.00	-.06	.39
Computers are out to get me	.11	-.13	.32
I have never been good at mathematics	.01	.05	-.09
I did badly at mathematics at school	-.01	-.11	.06
I slip into a coma whenever I see an equation	.08	.02	.09
Eigenvalues	7.29	1.74	1.32
% of variance	31.70	7.56	5.73
α	.82	.57	.82

Note: Factor loadings over .40 are in bold.

note: Factor loadings over .40 appear in bold

SPSS-anxiety therapy in that month). Another way to look at reliability is to see if people who are the same in terms of the construct being measured should get similar scores. So, if we took two people who were equally SPSS-phobic, then they should get equally high scores on the SAQ. Likewise, if we took two people who loved SPSS, they should both get equally low scores. It should be apparent that the SAQ would be an accurate measure of SPSS anxiety if we took someone who loved SPSS and so



LABCOAT LENI'S REAL RESEARCH 17.1

Worldwide addiction? ②

NICHOLS, L. A., & NICKI, R. (2004). *PSYCHOLOGY OF ADDICTIVE BEHAVIORS*, 185, 381–384.

In 2007 it was estimated that around 179 million people worldwide used the Internet. From the increasing popularity (and usefulness) of the Internet has emerged a serious and recognized problem of internet addiction. To research this construct it's helpful to be able to measure it, so Laura Nichols and Richard Nicki developed the Internet Addiction Scale (Nichols & Nicki, 2004). Nichols and Nicki's 36-item questionnaire contains items such as 'I have stayed on the Internet longer than I intended to' and 'My grades/work have suffered because of my Internet use' to which responses are made on a 5-point scale (Never, Rarely, Sometimes, Frequently, Always). (Incidentally, while researching this topic I encountered

an Internet addiction recovery whole host of resources (e.g. support groups, videos, podcasts) online for ages. It struck me a heroin addiction recovery center of free heroin in the reception area.

The data from 207 people is available in the file **Nichols & Nicki (2004).sav**. I dropped three items because they had low reliability compared to the other items. They performed a principal component analysis on the remaining items and found that the first component explains 60% of the variance. I want you to run some descriptive statistics on the data, which two items were dropped, inspect the variances, then inspect a correlation matrix of the three items that were dropped. Finally, he wants you to run a principal component analysis on the data. Answers are available in the companion website article).

was terrified of it and they got the same score! In statistical terms, test-retest reliability is based on the idea that individual items (or sets of items) should be consistent with the overall questionnaire. So, if we take someone's score on the overall score on the SAQ will be high; if the SAQ is reliable then the scores on some items from it the person's score on those items should also be high.

The simplest way to do this in practice is to use **split-half reliability**. This involves splitting the scale set into two randomly selected sets of items. A score for each item is calculated on each half of the scale. If a scale is reliable a person's score should be the same (or similar) to their score on the other half. Across many people, the scores from the two halves of the questionnaire should correlate very highly. The statistic computed between the two halves is the statistic computed in the split-half method, which is a sign of reliability. The problem with this method is that it depends on the way in which a set of data can be randomly split into two and so the results can vary depending on the way in which the data were split. To overcome this problem, we use a measure that is loosely equivalent to creating two sets of data by randomly splitting the data and computing the correlation coefficient for each split. The measure is equivalent to **Cronbach's alpha**, α , which is the most common measure of reliability.

$$\alpha = \frac{N^2 \overline{\text{cov}}}{\sum \sigma^2}$$

$$\sum s_{\text{item}}^2 + \sum \text{COV}_{\text{item}}$$

¹⁰ Although this is the easiest way to conceptualize Cronbach's α , whether or not it is the best of all possible split-half reliabilities depends on exactly how you calculate the split-half reliability (for computational details). If you use the Spearman–Brown formula, which takes item standard deviations into account, then Cronbach's α will be equal to the average split-half reliability only when all item standard deviations are equal; otherwise α will be smaller than the average. However, if you use a formula for reliability that does account for item standard deviations (such as Flanagan, 1937; Rulifson, 1947), then Cronbach's α will equal the average split-half reliability (see Cortina, 1993).

CHAPTER 17 EXPLORATORY FACTOR ANALYSIS

This equation may look complicated, but actually isn't. For each item on our scale, we calculate two things: the variance within the item, and the covariance between that item and any other item on the scale. Put another way, we can construct a variance-covariance matrix of all items. In this matrix the diagonal elements will be the variances for each particular item, and the off-diagonal elements will be covariances between pairs of items. The top half of the equation is simply the number of items (N) squared multiplied by the average covariance between items (the average of the off-diagonal elements in the mentioned variance-covariance matrix). The bottom half is the sum of all the item variances and item covariances (i.e., the sum of everything in the variance-covariance matrix).

There is a standardized version of the coefficient too, which essentially uses the same equation except that correlations are used rather than covariances, and the bottom half of the equation uses the sum of the elements in the correlation matrix of items (including the 1s that appear on the diagonal of that matrix). The normal alpha is appropriate for use when items on a scale are summed to produce a single score for that scale (the standardization is not appropriate in these cases). The standardized alpha is useful, though, when items on a scale are standardized before being summed.

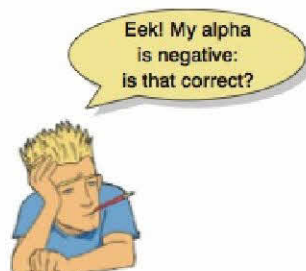
17.9.2. Interpreting Cronbach's α (some cautionary tales)

You'll often see in books or journal articles, or be told by people, that a value of .8 is an acceptable value for Cronbach's α ; values substantially lower indicate a scale is unreliable. Kline (1999) notes that although the generally accepted value of .8 is for cognitive tests such as intelligence tests, for ability tests a cut-off point of .7 is suitable. He goes on to say that when dealing with psychological constructs, even .7 can, realistically, be expected because of the diversity of the constructs involved. Some even suggest that in the early stages of research, values as low as .5 are acceptable (Nunnally, 1978). However, there are many reasons not to use these general guidelines, the least of which is that they distract you from thinking about what the value means in the context of the research you're doing (Pedhazur & Schmelkin, 1991).

We'll now look at some issues in interpreting alpha, which have been discussed particularly well by Cortina (1993) and Pedhazur and Schmelkin (1991). First, the value of alpha depends on the number of items on the scale. You'll notice that the top half of the equation for α includes the number of items squared. Therefore, as the number of items increases, α will increase. As such, it's possible to get a large value of α because of a large number of items on the scale, and not because your scale is reliable. For example, Cortina reports data from two scales, both of which have $\alpha = .8$. The first scale has only 5 items and the average correlation between items was a respectable .57; however, the second scale had 10 items with an average correlation between these items of a less respectable .35. Clearly the internal consistency of these scales differs, but according to Cronbach's alpha they are both equally reliable.

Second, people tend to think that alpha measures 'unidimensionality', or that the scale measures one underlying factor or construct. This is true when the scale is unidimensional (Cortina, 1993), but Cronbach (1951) demonstrated

factor underlying the data (see Cortina, 1993), but Grayson (2004) demonstrates that sets with the same α can have very different factor structures. He showed that the same α can be achieved in a scale with one underlying factor, with two moderately correlated factors, and with two uncorrelated factors. Cortina (1993) has also shown that with many items, and fairly high correlations between items ($r > .5$), α can reach values above .7 (.65 to .84). These results show that α should not be used as a measure of dimensionality'. Indeed, Cronbach (1951) suggested that if several factors exist, the formula should be applied separately to items relating to different factors. In conclusion, if your questionnaire has subscales, α should be applied separately to these subscales.



The final warning is about items that have a reverse phrasing. In the SAQ there is one item (question 3) that is phrased the opposite way around to all other items. The item was 'statistics make me nervous'. Compare this to any other item and you'll see it's a different response. For example, item 1 is 'statistics make me nervous then you'll strongly agree with this statement of 5 on our scale. For item 3, if you hate statistics then you're unlikely to excite you so you'll strongly disagree at the scale. These reverse-phrased items are important because of bias; participants will need to pay attention to the questions. For reverse phrasing doesn't matter; all that happens is you get a negative value for any reversed items (in fact, you'll see that item 3 has a negative factor loading of 17.9). However, these reverse-scored items will affect alpha. To see this, look at the equation for Cronbach's α . The top half incorporates the average squared deviation of items. If an item is reverse-phrased then it will have a negative correlation with the other items, hence the covariances between this item and other items. Since the average covariance is the sum of covariances divided by the number of items, by including a bunch of negative values we reduce the sum of covariances, and so also reduce Cronbach's α , because the top half of the equation gives the sum of squares. In some cases, it is even possible to get a negative value for Cronbach's α , since the magnitude of negative covariances is bigger than the magnitude of positive covariances. Cronbach's α doesn't make much sense, but it does happen, and whether you included any reverse-phrased items.

If you have reverse-phrased items then you also have to reverse the scoring of these items before you conduct reliability analysis. This is quite straightforward. In our data, we have one item which is currently scored as 1 = strongly disagree, 2 = disagree, 3 = neither, 4 = agree and 5 = strongly agree. This is fine for item 1 (statistics anxiety), but for item 3 (statistics make me nervous), disagreement indicates statistics anxiety. To reflect this, we need to reverse the scale such that 1 = strongly agree, 2 = agree, 3 = neither, 4 = disagree and 5 = strongly disagree. In doing so, an anxious person still gets a high score (they'd strongly disagree with it).

To reverse the scoring find the maximum value of your response scale and add 1 to it (so you get 6 in this case). Then for each person, subtract from it the score they actually got. Therefore, someone who originally scored 5 now scores $6 - 5 = 1$, and someone who scored 1 originally now gets $6 - 1 = 5$ in the middle of the scale with a score of 3 will still get $6 - 3 = 3$. It's a long time to do this for each person, but we can get SPSS to do it for us.



SELF-TEST Using what you learnt in Chapter 1, use the `compute` command to reverse-score item 3. (Check that you are simply changing the variable to 6 minus the original value.)



17.9.3. Reliability analysis in SPSS ②

Let's test the reliability of the SAQ using the data in SAQ.sav. You scored item 3 (see above), but if you can't be bothered then load

CHAPTER 17 EXPLORATORY FACTOR ANALYSIS

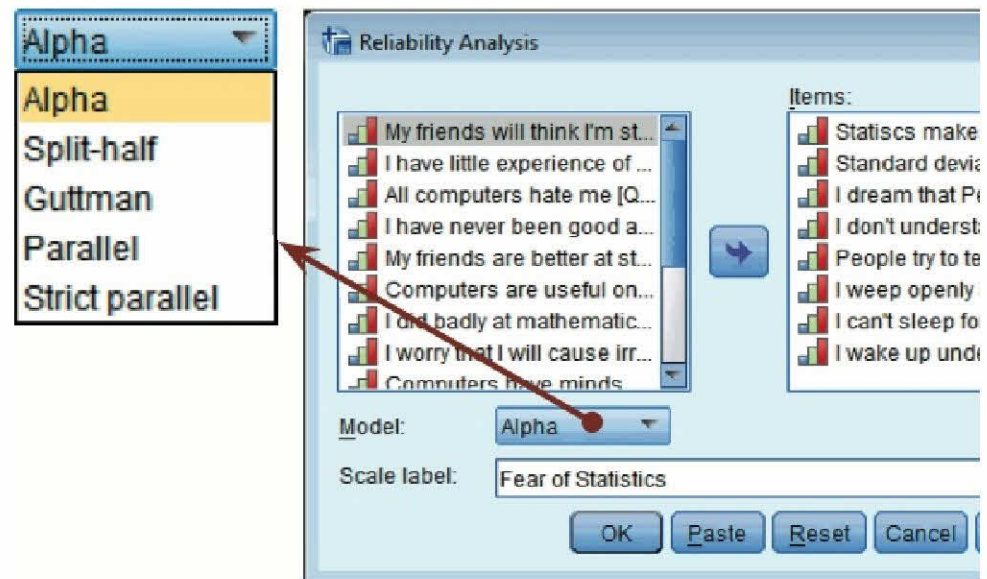


FIGURE 17.15 Main dialog box for reliability analysis.

Reversed).sav instead. Remember also that I said we should conduct reliability any subscales individually. If we use the results from our oblique rotation (O) then we have four subscales:

- 1 Subscale 1 (*Fear of statistics*): items 1, 3, 4, 5, 12, 16, 20, 21
- 2 Subscale 2 (*Peer evaluation*): items 2, 9, 19, 22, 23
- 3 Subscale 3 (*Fear of computers*): items 6, 7, 10, 13, 14, 15, 18
- 4 Subscale 4 (*Fear of mathematics*): items 8, 11, 17

To conduct each reliability analysis on these data you need to select the **Scale** menu item and click on **Reliability Analysis...** to display the dialog box in Figure 17.15. Drag any items from the list that you want to analyse (to begin with, let's do the items for the fear of statistics subscale: items 1, 3, 4, 5, 12, 16, 20 and 21) on the left-hand side of the dialog box and drag them to the box labelled **Items** (or click on the **➡** button). Remember that you can select several items at the same time if you hold down the **Ctrl** (**Cmd** on Mac) key while you select the variables.

There are several reliability analyses you can run, but the default option is α . You can change the method (e.g., to the split-half method) by clicking on the **Model** dropdown menu to reveal a drop-down list of possibilities, but the default method is a good one to use. It's a good idea to type the name of the scale (in this case 'Fear of Statistics') in the box labelled *Scale label* because this will add a header to the SPSS output with w

type in this box: typing a sensible name here will make your output easier to find.

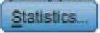
If you click on  you can access the dialog box in Figure 17.16. In this dialog box you can select several things, but the one most important for questionnaire reliability is: *Scale if item deleted*. This option tells us what the value of α would be if an item were deleted. If our questionnaire is reliable then we would not expect a single item to greatly affect the overall reliability. In other words, no item should cause a large decrease in α . If it does then you should consider dropping that item from the questionnaire to improve reliability.

FIGURE 17.16
Statistics
for reliability
analysis

Reliability Analysis: Statistics

Descriptives for

- ☐ Item
- ☐ Scale
- ☒ Scale if item deleted

Inter-Item

- ☒ Correlations
- ☐ Covariances

Summaries

- ☒ Means
- ☒ Variances
- ☒ Covariances
- ☒ Correlations

ANOVA Table

- ☒ None
- ☐ F test
- ☐ Friedman chi-square
- ☐ Cochran chi-square

☐ Hotelling's T-square

☐ Intraclass correlation coefficient

Model: Two-Way Mixed Type: Consistency

Confidence interval: 95 % Test value: 0

Continue Cancel Help

The inter-item correlations and covariances (and summaries) provide information about the relationship between items on our scale. We should also look at the item means and variances from our factor analysis, so there is little point in selecting these options. The *F test*, *Friedman chi-square* (if your data are ranked), *Cochran chi-square* (if your data are dichotomous), and *Hotelling's T-square* use these tests to compare different items on the questionnaire. These tests might be useful to check for similar distributional properties (i.e., the same average value), but for the purposes of factor analysis, they will inevitably give misleading results even when only small differences exist between the questionnaires.

You can also request an **intraclass correlation coefficient (ICC)**. The ICCs that we encountered earlier in this book measure the relationship between variables that measure different things. For example, the correlation between Omega and Satanism involves two classes of measures: the type of religious belief and their religious beliefs. Intraclass correlations measure the relationship between variables that measure the same thing (i.e., variables within the same class). They are used in comparing paired data (such as twins) on the same measure, or in assessing consistency between judges' ratings of a set of objects (hence the name 'intraclass'). (See the reliability statistics in SPSS). If you'd like to know more, see Section 17.3.

Use the simple set of options in Figure 17.16 to run a basic reliability analysis. Click on the **Continue** button to return to the main dialog box and then click on the **OK** button to

17.9.4. Reliability analysis output ②

Output 17.13 shows the results of this basic reliability analysis for the subscale. The value of Cronbach's α is presented in a small table as a measure of the reliability of the scale. Bearing in mind what we've already noted about the number of items, and how daft it is to apply general rules, we're looking for a value in the region of about .7 to .8. In this case α is .821, which is certainly in the region of good reliability (Kline (1999)), and probably indicates good reliability.

CHAPTER 17 EXPLORATORY FACTOR ANALYSIS

Item-Total Statistics			
	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation
Statistics makes me cry	21.76	21.442	.821
Standard deviations excite me	20.72	19.825	.823
I dream that Pearson is attacking me with correlation coefficients	21.35	20.410	.823
I don't understand statistics	21.41	20.942	.823
People try to tell you that SPSS makes statistics easier to understand but it doesn't	20.97	20.639	.823
I weep openly at the mention of central tendency	21.25	20.451	.823
I can't sleep for thoughts of eigenvectors	20.51	21.176	.823
I wake up under my duvet thinking that I am trapped under a normal distribution	20.96	19.939	.823

Reliability Statistics

Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
.821	.823	8

OUTPUT 17.13

In the table labelled *Item-Total Statistics* the column labelled *Corrected Correlation* has the correlations between each item and the total score from the scale. In a reliable scale all items should correlate with the total. So, we're looking for items that don't correlate with the overall score from the scale: if any of these values is below about .3 then we've got problems, because it means that a particular item does not correlate very well with the scale overall. Items with low correlations may have to be deleted. For these data, all data have item-total correlations above .3, which is encouraging.

The values in the column labelled *Cronbach's Alpha if Item Deleted* are the values of the overall α if that item isn't included in the calculation. As such, they reflect the Cronbach's α that would be seen if a particular item were deleted. The overall Cronbach's α is .821, and so all values in this column should be around that same value. We're actually looking for values of alpha greater than the overall α . If you think about it, if the deletion of an item increases Cronbach's α then this means that the deletion of that item improves the reliability. Therefore, any items that have values of α in this column greater than the overall α may need to be deleted from the scale to improve its reliability. None of the items would increase alpha if they were deleted, which is good news. It's worth noting that none of the items do need to be removed at this stage then you should rerun your factor analysis to make sure that the deletion of the item has not affected the factor structure.



SELF-TEST

Run reliability analyses on the other three subscales.



Just to illustrate the importance of reverse-scoring items before running reliability analysis, Output 17.14 shows the reliability analysis for the fear of statistics subscale on the original data (i.e., without item 3 being reverse-scored). Note that the reliability is considerably lower (.605 rather than .821). Also, note that this item has a negative item-total correlation (which is a good way to spot if you have a potential reverse-scored item). Finally, note that for item 3, the α if it is deleted is higher than the overall α .

Item-Total Statistics			
	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation
Statistics makes me cry	20.93	12.125	.505
Standard deviations excite me	20.72	19.825	-.549
I dream that Pearson is attacking me with correlation coefficients	20.52	11.447	.526
I don't understand statistics	20.58	11.714	.466
People try to tell you that SPSS makes statistics easier to understand but it doesn't	20.14	11.739	.501
I weep openly at the mention of central tendency	20.42	11.584	.529
I can't sleep for thoughts of eigenvectors	19.68	12.107	.353
I wake up under my duvet thinking that I am trapped under a normal distribution	20.13	11.189	.541

Reliability Statistics		
Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
.605	.641	8

OUTPUT 17.14

is .8. That is, if this item were deleted then the reliability would im about .8. This, I hope, illustrates that failing to reverse-score items oppositely to other items on the scale will mess up your reliability :

Let's now look at our subscale of peer evaluation. For our subs you should get the output in Output 17.15. The overall reliability ing to bake a cake for. The overall α is quite low, and although what Kline says we should expect for this kind of social science d: statistics subscale and (as we shall see) the other two. The scale ha to seven, eight and three on the other scales, so its reliability relati not going to be dramatically affected by the number of items. Th labelled *Corrected Item-Total Correlation* are all around .3, and sm. results again indicate questionable internal consistency and identif problem. The values in the column labelled *Cronbach's Alpha if Ite* none of the items here would increase the reliability if they were de in this column are less than the overall reliability of .57. The item: quite diverse themes of peer evaluation, and this might explain the ency; we probably need to rethink this subscale.

Moving on to the fear of computers subscale, Output 17.16 shov which is pretty good. The values in the column labelled *Corrected*

OUTPUT 17.15

Item-Total Statistics			
	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation
My friends will think I'm stupid for not being able to cope with SPSS	11.46	8.119	.339

My friends are better at statistics than me	10.24	6.395	.391
Everybody looks at me when I use SPSS	10.79	7.381	.316
My friends are better at SPSS than I am	10.20	7.282	.378
If I'm good at statistics my friends will think I'm a nerd	9.65	7.988	.239

Reliability Statistics

Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
.570	.572	5

CHAPTER 17 EXPLORATORY FACTOR ANALYSIS

Item-Total Statistics					
	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Squared Multiple Correlation	Cronbach's Alpha if Deleted
I have little experience of computers	15.87	17.614	.619	.398	
All computers hate me	15.17	17.737	.619	.395	
Computers are useful only for playing games	15.81	20.736	.400	.167	
I worry that I will cause irreparable damage because of my incompetence with computers	15.64	18.809	.607	.384	
Computers have minds of their own and deliberately go wrong whenever I use them	15.22	18.719	.577	.350	
Computers are out to get me	15.33	19.322	.491	.250	
SPSS always crashes when I try to use it	15.52	17.832	.647	.447	

Reliability Statistics		
Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
.823	.821	7

are again all above .3, which is also good. The values in the column labelled *Alpha if Item Deleted* show that none of the items would increase the reliability deleted. This indicates that all items are positively contributing to the overall reliability.

Finally, for the fear of maths subscale, Output 17.17 shows an overall reliability which indicates good reliability. The values in the column labelled *Corrected Correlation* are all above .3, which is good, and the values in the column labelled *Alpha if Item Deleted* indicate that none of the items here would increase the reliability were deleted because all values in this column are less than the overall reliability.

Item-Total Statistics					
	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Squared Multiple Correlation	Cronbach's Alpha if Deleted
I have never been good at mathematics	4.72	2.470	.684	.470	
I did badly at mathematics at school	4.70	2.453	.682	.467	
I slip into a coma whenever I see an equation	4.49	2.504	.652	.425	

Reliability Statistics		
Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
.819	.819	3



CRAMMING SAM'S TIPS

Reliability

- Reliability analysis is used to measure the consistency of a measure.

- Remember to reverse-score any items that were reverse-phrased on the original questionnaire analysis.
- Run separate reliability analyses for all subscales of your questionnaire.
- Cronbach's α indicates the overall reliability of a questionnaire, and values around .8 are good (the like).
- The *Cronbach's Alpha if Item Deleted* column tells you whether removing an item will improve reliability. Values greater than the overall reliability indicate that removing that item will improve the overall reliability. Look for items that dramatically increase the value of α and remove them.
- If you remove items, rerun your factor analysis to check that the factor structure still holds.

17.10. How to report reliability analysis

You can report the reliabilities in the text using the symbol α and remember that Cronbach's α can't be larger than 1 (we drop the zero before the decimal following APA practice):

- ✓ The fear of computers, fear of statistics and fear of maths subscale had high reliabilities, all Cronbach's $\alpha = .82$. However, the fear of change subscale had relatively low reliability, Cronbach's $\alpha = .5$

However, the most common way to report reliability analysis with a table is to report the values of Cronbach's α as part of the table. For example, in Table 17.1 notice that in the last row of the table we report Cronbach's α for each subscale in turn.

17.11. Brian's attempt to woo Jane ①

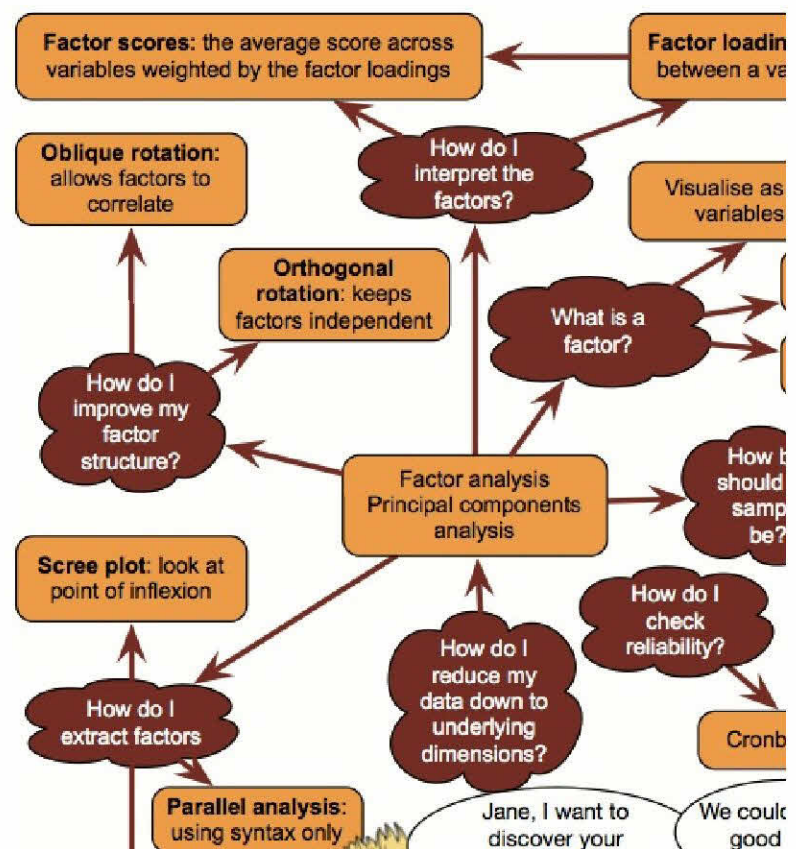




FIGURE 17.17
What Brian learnt from this chapter

CHAPTER 17 EXPLORATORY FACTOR ANALYSIS

17.12. What next? ②

At the age of 23 I took it upon myself to become a living homage to the digests I furiously devoured articles and books on statistics (some of them I even un- mentally chewed over them, I broke them down with the stomach acid of my stomach, I stripped them of their goodness and nutrients, I compacted them down, and two years I forced the smelly brown remnants of those intellectual meals out of me in the form of a book. I was mentally exhausted at the end of it. ‘It’s a good job I’ll not do that again’, I thought.

17.13. Key terms that I’ve discovered

Alpha factoring	Factor matrix	Principal component analysis (PCA)
Anderson–Rubin method	Factor scores	Promax
Common factor	Factor transformation matrix, Λ	Quartimax
Common variance	Intraclass correlation coefficient (ICC)	Random variance
Communality	Kaiser’s criterion	Rotation
Component matrix	Latent variable	Scree plot
Confirmatory factor analysis	Kaiser–Meyer–Olkin (KMO)	Singularity
Cronbach’s α	measure of sampling adequacy	Split-half reliability
Direct oblimin	Oblique rotation	Structure matrix
Extraction	Orthogonal rotation	Unique factor
Equamax	Pattern matrix	Unique variance
Factor analysis		Varimax
Factor loading		

17.14. Smart Alex’s tasks

- **Task 1:** Rerun the analysis in this chapter using principal component analysis and compare the results to those in the chapter. (Set the iterations to convergence to 0.0001.)
- **Task 2:** The University of Sussex constantly seeks to employ the best people as statistics lecturers. They wanted to revise the ‘Teaching of Statistics for Scientific Experiments’ (TOSSE) questionnaire, which is based on Bland’s theory that says that good statistics lecturers should have: (1) a profound love of statistics; (2) an understanding of experimental design; (3) a love of teaching; and (4) a complete absence of interpersonal skills. These characteristics should be related (i.e., correlated). The University revised this questionnaire to become the ‘Teaching of Statistics for Scientific Experiments – Revised’ (TOSSE-R). They gave this questionnaire to 20 statistics lecturers around the world to see if it supported Bland’s theory.

tionnaire is in Figure 17.18, and the data are in **TOSSE-R.sav**. Conduct a factor analysis (with appropriate rotation) and interpret the factor structure. ②

- **Task 3:** Dr Sian Williams (University of Brighton) devised a questionnaire on organizational ability. She predicted five factors to do with organizational ability: (1) preference for organization; (2) goal achievement; (3) planning approach; (4) tolerance of delays; and (5) preference for routine. These dimensions are theoretically independent. Williams' questionnaire contains 28 items using a 7-point scale (1 = strongly disagree, 4 = neither, 7 = strongly agree). She gave it to 100 participants. Run a principal component analysis on the data in **Williams.sav**. ②

- **Task 4:** Zibarras, Port, and Woods (2008) looked at the relationship between personality and creativity. They used the Hogan Development Surveys 11 dysfunctional dispositions of employed adults: being cautious, detached, passive-aggressive, arrogant, manipulative, perfectionist, and dependent. Zibarras et al. wanted to reduce these 11 dispositions based on parallel analysis, found that they could be reduced to 5 factors. They ran a principal component analysis with varimax rotation (Zibarras et al. (2008).sav) to see which personality dimensions were related to creativity (see page 210 of the original paper). ②

Answers can be found on the companion website.

FIGURE 17.18
The TOSSE-R
questionnaire

Teaching of Statistics for Scientific Experiments – Revised (TOSSE-R)		SD
1.	I once woke up in a vegetable patch hugging a turnip that I'd mistakenly dug up thinking it was Roy's largest root	☐
2.	If I had a big gun I'd shoot all the students I have to teach	☐
3.	I memorize probability values for the F -distribution	☐
4.	I worship at the shrine of Pearson	☐
5.	I still live with my mother and have little personal hygiene	☐
6.	Teaching others makes me want to swallow a large bottle of bleach because the pain of my burning oesophagus would be light relief in comparison	☐
7.	Helping others to understand sums of squares is a great feeling	☐
8.	I like control conditions	☐
9.	I calculate 3 ANOVAs in my head before getting out of bed	☐
10.	I could spend all day explaining statistics to people	☐
11.	I like it when I've helped people to understand factor rotation	☐
12.	People fall asleep as soon as I open my mouth to speak	☐
13.	Designing experiments is fun	☐
14.	I'd rather think about appropriate dependent variables than go to the pub	☐
15.	I soil my pants with excitement at the mere mention of Factor Analysis	☐
16.	Thinking about whether to use repeated- or independent-measures thrills me	☐
17.	I enjoy sitting in the park contemplating whether to use participant observation in my next experiment	☐
18.	Standing in front of 300 people in no way makes me lose control of my bowels	☐
19.	I like to help students	☐
20.	Passing on knowledge is the greatest gift you can bestow on an individual	☐
21.	Thinking about Bonferroni corrections gives me a tingly feeling in my groin	☐
22.	I quiver with excitement when thinking about designing my next	☐

experiment

- 22. I often spend my spare time talking to the pigeons ... and even they die of boredom ☐
- 23. I tried to build myself a time machine so that I could go back to the 1930s and follow Fisher around on my hands and knees licking the floor on which he'd just trodden ☐
- 25. I love teaching ☐
- 26. I spend lots of time helping students ☐
- 27. I love teaching because students have to pretend to like me or they'll get bad marks ☐
- 28. My cat is my only friend ☐

CHAPTER 17 EXPLORATORY FACTOR ANALYSIS

17.15. Further reading

- Cortina, J. M. (1993). What is coefficient alpha? An examination of theory and application. *Journal of Applied Psychology*, 78, 98–104. (A very readable paper on Cronbach's α .)
- Dunteman, G. E. (1989). *Principal components analysis*. Sage University Paper Series on Applications in the Social Sciences, 07-069. Newbury Park, CA: Sage. (This monograph is at a high level but comprehensive.)
- Pedhazur, E., & Schmelkin, L. (1991). *Measurement, design and analysis*. Hillsdale, NJ: Lawrence Erlbaum. (Chapter 22 is an excellent introduction to the theory of factor analysis.)
- Tabachnick, B. G., & Fidell, L. S. (2012). *Using multivariate statistics* (6th ed.). Boston: Allyn and Bacon.

18

Categorical data



FIGURE 18.1

Midway through writing the second edition of this book, things became a little strange



18.1. What will this chapter tell me? (

We discovered in the previous chapter that I wrote a book. This book tells good things about writing books. The main benefit is that your family will like them. They're not *that* impressed, because they think that a good book is like *Harry Potter* and that people should queue outside bookshops for the next instalment of *Discovering statistics* My parents are, consequently, confused about how this book is seen as successful, yet I don't get invited to write more. Nevertheless, given that my family don't really understand what I do, I prove that I do *something*. The size of this book and the fact it has

added bonus because it makes me look cleverer than I actually am. The price to pay, which is immeasurable mental anguish. In England we have emotions, because we fear that if they get out into the open, civilization will collapse, so I definitely will not mention that the writing process was so stressful that I came within one of Fuzzy's whiskers of a time or two years to recover, just in time to start thinking about the time it was worth it because the feedback suggests that some people found the book helpful. Of course, the publishers focus less on the book's helpfulness and more on comparisons with other books. They have databases that have sales

CHAPTER 18 CATEGORICAL DATA

and its competitors in different ‘markets’ (you are not a person, you are a ‘company’, you don’t live in a country, you live in a ‘market’), and they gibber and twitch: soles creating pink frequency distributions (with 3-D effects) of these values. They get are frequency data (the number of books sold in a certain period of time) if they wanted to compare sales of this book to its competitors, in different countries they would need to read this chapter because it’s all about analysing data for which only the frequency with which events occur. Of course, they won’t read this chapter, they should ...

18.2. Analysing categorical data ①

So far we have looked at fitting models with categorical predictor variables, predicting a continuous outcome variable. Sometimes, however, we want to predict categorical outcome variables. In other words, we want to predict into which category an entity falls. For example, we might want to predict whether someone is male or female, whether a person is married or not, for which political party a person voted, whether a tumour is benign or malignant, whether a sports team will win, lose or draw. In all of these cases, an entity can only belong to one category, for example a woman can be pregnant or not; she can’t be both. The next two chapters deal with statistical models for categorical outcomes. We begin with some basic models of associations between categorical variables, then in the next chapter we’ll move on to look at predicting categorical outcomes from both categorical and continuous predictor variables.

18.3. Theory of analysing categorical data ①

We will begin by looking at the simplest situation that you could encounter; that is, analysing the relationship between two categorical variables. With categorical variables, you can’t use the mean or any similar statistic because the mean of a categorical variable is completely meaningless: the numeric values you attach to different categories are arbitrary, and the mean of those numeric values will depend on how many members each category has. Therefore, when we’ve measured only categorical variables, we analyse the data in terms of things that fall into each combination of categories (i.e., the frequencies). For example, a researcher was interested in whether animals could be trained to line-dance. He had 200 cats and tried to train them to line-dance by giving them either food or affection as a reward for dance-like behaviour. At the end of the week he counted how many cats could line-dance and how many could not. There are two categorical variables here: **training** (the animal was trained using either food or affection, not both) and **dance** (the animal either learnt to line-dance or it did not). By combining categories, we end up with four different categories. All we then need to do is to count how many cats fall into each category. We can tabulate these frequencies as in Table 18.1, which shows the data for the

we can tabulate these frequencies as in Table 18.1 (which shows the data for the number of cats that line-dance relate to the type of training used?) and this is known as a **contingency table**.

18.3.1. Pearson's chi-square test ①

If we want to see whether there's a relationship between two categorical variables (the number of cats that line-dance relate to the type of training used?) we can use a **chi-square test** (Fisher, 1922; Pearson, 1900). This is an extremely elegant sta-

TABLE 18.1 Contingency table showing how many cats will line-dance different rewards

		<i>Training</i>	
		<i>Food as Reward</i>	<i>Affection as Reward</i>
Could They Dance?	Yes	28	4
	No	10	11
	Total	38	16

on the simple idea of comparing the frequencies you observe in categories with the frequencies you might expect to get in those categories by chance (equation (2.6)) that if we want to calculate the fit (or total error) of a model, we square the differences between the observed values of the outcome, and the values that come from the model:

$$\text{Total error} = \sum_{i=1}^n (\text{observed}_i - \text{model}_i)^2$$

This equation was the basis of our sums of squares in regression analysis. When we have categorical data we use essentially the same equation. There is a slight modification: we divide by the model scores as well, which is actually much the same as dividing the sum of squares by the degrees of freedom in ANOVA. Basically, we are calculating the standardized deviation for each observation. If we add all of these standardized deviations together, the resulting statistic is Pearson's chi-square (χ^2) given by:

$$\chi^2 = \sum \frac{(\text{observed}_{ij} - \text{model}_{ij})^2}{\text{model}_{ij}}$$

in which i represents the rows in the contingency table and j represents the columns. The observed data are, obviously, the frequencies in Table 18.1, but what the model is. When we have categorical predictors but a continuous outcome variable (ANOVA) the model we use is group means, but, as I've mentioned, when we have a categorical outcome variable so we work with expected frequencies. Therefore, we use 'expected frequencies'. One way to estimate the expected frequencies would be to say 'well, we've got 200 cats in total, and four categories, so the expected value is simply $200/4 = 50$ '. This would be fine if, for example, we had 200 cats that had affection as a reward as we did cats that had food as a reward. 38 got food and 162 got affection as a reward. Likewise there are 10 cats that could and couldn't dance. To take account of these inequalities we use the expected frequencies for each cell in the table (in this case there are

expected frequencies for each cell in the table (in this case there are column and row totals for a particular cell to calculate the expected

$$\text{model}_{ij} = E_{ij} = \frac{\text{row total}_i \times \text{column total}_j}{n}$$

In which n is the total number of observations (in this case 200).
 expected frequencies for the four cells within our table (row totals abbreviated to RT and CT, respectively):

CHAPTER 18 CATEGORICAL DATA

$$\begin{aligned}\text{model}_{\text{Food, Yes}} &= \frac{RT_{\text{Yes}} \times CT_{\text{Food}}}{n} = \frac{76 \times 38}{200} = 14.44 \\ \text{model}_{\text{Food, No}} &= \frac{RT_{\text{No}} \times CT_{\text{Food}}}{n} = \frac{124 \times 38}{200} = 23.56 \\ \text{model}_{\text{Affection, Yes}} &= \frac{RT_{\text{Yes}} \times CT_{\text{Affection}}}{n} = \frac{76 \times 162}{200} = 61.56 \\ \text{model}_{\text{Affection, No}} &= \frac{RT_{\text{No}} \times CT_{\text{Affection}}}{n} = \frac{124 \times 162}{200} = 100.44\end{aligned}$$

Given that we now have these model values, all we need to do is take each value of our data table, subtract from it the corresponding model value, square it, and then divide by the corresponding model value. Once we've done this for each cell in the table, we just add them up!

$$\begin{aligned}\chi^2 &= \frac{(28 - 14.44)^2}{14.44} + \frac{(10 - 23.56)^2}{23.56} + \frac{(48 - 61.56)^2}{61.56} + \frac{(114 - 100.44)^2}{100.44} \\ &= \frac{13.56^2}{14.44} + \frac{-13.56^2}{23.56} + \frac{-13.56^2}{61.56} + \frac{13.56^2}{100.44} \\ &= 12.73 + 7.80 + 2.99 + 1.83 \\ &= 25.35\end{aligned}$$

This statistic can then be checked against a distribution with known properties: the **chi-square distribution**. All we need to know is the degrees of freedom, which are calculated as $(r-1)(c-1)$ in which r is the number of rows and c is the number of columns. Another way to think of it is as the number of levels of each variable minus one, multiplied together. In this case we get $df = (2 - 1)(2 - 1) = 1$. If you were doing this by hand, you would find a critical value for the chi-square distribution with $df = 1$. If your observed value was bigger than this critical value you would say that there was a significant relationship between the two variables. These critical values are produced in tables, and for $df = 1$ they are 3.84 ($p = .05$) and 6.63 ($p = .01$). So because the observed chi-square is bigger than these values it is significant at $p < .01$. However, SPSS will give you the precise probability of obtaining a chi-square statistic at least as big as the observed value (25.35) if there were no association between the variables in the population.

18.3.2. Fisher's exact test ①

There is one problem with the chi-square test, which is that the sampling distribution of the test statistic has an *approximate* chi-square distribution. The larger the sample size, the better this approximation becomes, and in large samples the approximation is good enough that we do not worry about the fact that it is an approximation. However, in small samples the approximation is not good enough, making significance tests of the chi-square test inaccurate. This is why you often read that to use the chi-square test the expected frequencies in each cell must be greater than 5 (see Section 18.4). When the expected frequencies are greater than 5, the sampling distribution is probably close enough to a chi-square distribution for us not to worry. However, when the expected frequencies

low, it probably means that the sample size is too small and that the value of the test statistic is too deviant from a chi-square distribution to be reliable.

Fisher came up with a method for computing the exact probability that is accurate when sample sizes are small. This method is called Fisher's exact test (Fisher, 1922) even though it's not so much a test as a way of computing the probability of the chi-square statistic. This procedure is normally used for 2x2 tables (i.e., two variables each with two options) and with small sample sizes. It can be used on larger contingency tables and with large samples, but for large tables it becomes computationally intensive and you might find it difficult to give you an answer. In large samples there is really no point because you can overcome the problem of small samples, so you don't need to use it.

18.3.3. The likelihood ratio

An alternative to Pearson's chi-square is the likelihood ratio statistic, which is based on maximum-likelihood theory. The general idea behind this theory is to fit a model to the data and create a model for which the probability of obtaining the observed data is maximized, then you compare this model to the probability of obtaining the observed data under the null hypothesis. The resulting statistic is, therefore, based on the ratio of the observed frequencies with those predicted by the model:

$$L\chi^2 = 2 \sum \text{observed}_{ij} \ln \left(\frac{\text{observed}_{ij}}{\text{model}_{ij}} \right)$$

in which i and j are the rows and columns of the contingency table. The natural logarithm (this is the standard mathematical function that we can find on our calculator usually labelled as $\ln$ or $\log_e$) of the ratio of the observed values as in the previous section, this would give us:

$$\begin{aligned} L\chi^2 &= 2 \left[28 \times \ln \left(\frac{28}{14.44} \right) + 10 \times \ln \left(\frac{10}{23.56} \right) + 48 \times \ln \left(\frac{48}{61.56} \right) + 114 \times \ln \left(\frac{114}{235.6} \right) \right] \\ &= 2 [28 \times 0.662 + 10 \times -0.857 + 48 \times -0.249 + 114 \times 0.127] \\ &= 2 [18.54 - 8.57 - 11.94 + 14.44] \\ &= 24.94 \end{aligned}$$

As with Pearson's chi-square, this statistic has a chi-square distribution with 1 degree of freedom (in this case 1). As such, it is tested in the same way as the chi-square test. The critical value of chi-square for the number of degrees of freedom is 3.84. Before, the value we have here will be significant because it is bigger than 3.84.

of 3.84 ($p = .05$) and 6.63 ($p = .01$). For large samples this statistic is as Pearson's chi-square, but is preferred when samples are small.

18.3.4. Yates's correction ②

When you have a 2×2 contingency table (i.e., two categorical variables) then Pearson's chi-square tends to produce significance values (in other words, it tends to make a Type I error). Therefore, Yates

CHAPTER 18 CATEGORICAL DATA

to the Pearson formula (usually referred to as **Yates's continuity correction**). This is that when you calculate the deviation from the model (the observed_{ij} – model_{ij}) you subtract 0.5 from the absolute value of this deviation before it. In plain English this means you calculate the deviation, ignore whether it is negative, subtract 0.5 from the value and then square it. Pearson's equation then

$$\chi^2 = \sum \frac{(|\text{observed}_{ij} - \text{model}_{ij}| - 0.5)^2}{\text{model}_{ij}}$$

For the data in our example this just translates into:

$$\begin{aligned}\chi^2 &= \frac{(13.56 - 0.5)^2}{14.44} + \frac{(13.56 - 0.5)^2}{23.56} + \frac{(13.56 - 0.5)^2}{61.56} + \frac{(13.56 - 0.5)^2}{100.44} \\ &= 11.81 + 7.24 + 2.77 + 1.70 \\ &= 23.52\end{aligned}$$

The key thing to note is that it lowers the value of the chi-square statistic and makes it less significant. Although this seems like a nice solution to the problem, a fair bit of evidence that this overcorrects and produces chi-square values that are too low. Howell (2012) provides an excellent discussion of the problem with Yates's continuity correction if you're interested; all I will say is that although it's worth knowing, it's probably best ignored.

18.3.5. Other measures of association ①

There are measures of the strength of association other than the chi-square. These measures modify the chi-square statistic to take account of sample size and degrees of freedom and try to restrict the range of the test statistic from 0 to 1 (to make them comparable to the correlation coefficient described in Chapter 7). Three related measures are

- **Phi:** This statistic is accurate for 2 × 2 contingency tables. However, for more than two dimensions the value of phi may not lie between 0 and 1. A chi-square value can exceed the sample size. Therefore, Pearson suggested the contingency coefficient.
- **Contingency coefficient:** This coefficient ensures a value between 0 and 1. Unfortunately, it seldom reaches its upper limit of 1 and for this reason Cramér's V is an alternative denoted by V.
- **Cramér's V:** When both variables have only two categories, phi and Cramér's V are identical. However, when variables have more than two categories Cramér's V is the preferred measure.

can attain its maximum of 1 – unlike the other two – and so it is the most

18.3.6. Several categorical variables: loglinear analysis

So far we've looked at situations in which there are only two categorical variables. We often want to analyse more complex contingency tables in which there are three or more variables. For example, suppose we took the example we've just used but added a third variable, such as gender.

data from a sample of 70 dogs. We might want to compare the behaviour of cats. We would now have three variables: **Animal** (dog or cat), **Training** (affection as reward) and **Dance** (did they dance or not?). This could be analysed with Pearson chi-square and instead has to be analysed with a technique called multiple regression.

18.3.6.1. Chi-square as regression ④



To begin with, let's have a look at how our simple chi-square example can be viewed as a regression model. Although we already know about as much as we can about a chi-square test, if we want to understand more complex situations, it's a lot easier if we consider our model as a general linear model (i.e. the same as the general linear models we've considered in this book take the general form:

$$\text{outcome}_i = (\text{model}) + \text{error}_i$$

For example, when we encountered multiple regression in Chapter 8, the model was written as (see equation (8.6)):

$$Y_i = (b_0 + b_1X_{1i} + b_2X_{2i} + \dots + b_nX_{ni}) + \varepsilon_i$$

Also, when we came across one-way ANOVA, we adapted this regression model to analyze our Viagra example, as (see equation (11.1)):

$$\text{Libido}_i = b_0 + b_2\text{High}_i + b_1\text{Low}_i + \varepsilon_i$$

The *t*-test was conceptualized in a similar way. In all cases the same model is used, it's just the complexity of the model that changes. With categorical variables, the same model in much the same way as with regression to produce categorical outcomes. In the current example we have two categorical variables: training (food reward) and dance (yes they did dance or no they didn't dance). Both variables have two categories. We can represent each one with a single dummy variable (see Section 18.2). For training, the 'affection' category is coded as 0 and the other as 1. So for training, we could code the 'affection' variable as 1, and we could code the dancing variable as 1 for 'yes' and 0 for 'no' (see Table 18.2).

This situation might be familiar if you think back to factorial ANOVA in which we also had two variables as predictors. In that situation, where we have two variables, the general linear model became (think back to equation (8.6)):

$$\text{Outcome}_i = (b_0 + b_1A_i + b_2B_i + b_3AB_i) + \varepsilon_i$$

TABLE 18.2 Coding scheme for dancing cats

<i>Training</i>	<i>Dance</i>	<i>Dummy (Training)</i>	<i>Dummy (Dance)</i>	<i>Inter.</i>
Food	No	0	0	
Food	Yes	0	1	
Affection	No	1	0	
Affection	Yes	1	1	

CHAPTER 18 CATEGORICAL DATA

in which A represents the first variable, B represents the second and AB represents the interaction between the two variables. Therefore, we can construct a linear model using these dummy variables that is exactly the same as the one we used for factor ANOVA (above). The interaction term will be the training variable multiplied by the dance variable (look at Section 10.3.2, and if it doesn't make sense look back to Section 13.2). The coding is exactly the same as this example):

$$\text{outcome}_i = (\text{model}) + \text{error}_i$$

$$\text{outcome}_{ij} = (b_0 + b_1 \text{Training}_i + b_2 \text{Dance}_j + b_3 \text{Interaction}_{ij}) + \varepsilon_{ij}$$

However, because we're using categorical data, to make this model linear we actually use log values and so the actual model becomes:¹

$$\ln(O_i) = \ln(\text{model}) + \ln(\varepsilon_i)$$

$$\ln(O_{ij}) = (b_0 + b_1 \text{Training}_i + b_2 \text{Dance}_j + b_3 \text{Interaction}_{ij}) + \ln(\varepsilon_{ij})$$

The training and dance variables and the interaction can take the values 0 and 1 depending on which combination of categories we're looking at (Table 18.2). Therefore, to find out what the b -values represent in this model we can do the same as we did for ANOVA and look at what happens when we replace training and dance with 0 and 1. To begin with, let's see what happens when we look at when training and dance are both zero. This situation represents the category of cats that got food reward and didn't line-dance. When we used this sort of model for the t -test and ANOVA the observed outcomes were taken from the observed data: we used the group means (e.g., see Section 11.2.1). However, with a categorical outcome we use the observed frequencies rather than observed means). In Table 18.1 we saw that there were 10 cats that had food reward and didn't line-dance. If we use this as the observed outcome then this can be written as (if we ignore the error term for the time being):

$$\ln(O_{ij}) = b_0 + b_1 \text{Training}_i + b_2 \text{Dance}_j + b_3 \text{Interaction}_{ij}$$

For cats that had food reward and didn't dance, the training and dance variables and interaction will all be 0 and so the equation reduces to:

$$\ln(O_{\text{Food, No}}) = b_0 + (b_1 \times 0) + (b_2 \times 0) + (b_3 \times 0)$$

$$\ln(O_{\text{Food, No}}) = b_0$$

$$\ln(10) = b_0$$

¹ See Appendix A.

$$b_0 = 2.303$$

Therefore, b_0 in the model represents the log of the observed value when all of the other variables are zero. As such it's the log of the observed value of the base category (in this case, the category that got food and didn't dance).

¹ Actually, the convention is to denote b_0 as η and the b -values as β , but I think these notational changes would only serve to confuse people so I'm sticking with b because I want to emphasize the similarities to regression.

Now, let's see what happens when we look at cats that had affection and didn't dance. In this case, the training variable is 1 and the dance variable is still 0. Also, our outcome now changes to be the observed frequency for cats that received affection and didn't dance (from Table 18.1 we can see the value is 114). Therefore, the equation becomes:

$$\ln(O_{\text{Affection, No}}) = b_0 + (b_1 \times 1) + (b_2 \times 0) + (b_3 \times 0)$$

$$\ln(O_{\text{Affection, No}}) = b_0 + b_1$$

$$b_1 = \ln(O_{\text{Affection, No}}) - b_0$$

Remembering that b_0 is the expected value for cats that had food and didn't dance, we can get:

$$\begin{aligned} b_1 &= \ln(O_{\text{Affection, No}}) - \ln(O_{\text{Food, No}}) \\ &= \ln(114) - \ln(10) \\ &= 4.736 - 2.303 \\ &= 2.433 \end{aligned}$$

The important thing is that b_1 is the difference between the log of the observed frequency for cats that received affection and didn't dance, and the log of the expected frequency for cats that received food and didn't dance. Put another way, within the group that didn't dance it represents the difference between those trained using food and those trained using affection.

Now, let's see what happens when we look at cats that had food and danced. In this case, the training variable is 0, the dance variable is 1 and the outcome is 28. Our outcome now changes to be the observed frequency for cats that received food and danced (from Table 18.1 we can see the value is 28). Therefore, the equation becomes:

$$\ln(O_{\text{Food, Yes}}) = b_0 + (b_1 \times 0) + (b_2 \times 1) + (b_3 \times 0)$$

$$\ln(O_{\text{Food, Yes}}) = b_0 + b_2$$

$$b_2 = \ln(O_{\text{Food, Yes}}) - b_0$$

Remembering that b_0 is the expected value for cats that had food and didn't dance, we can get:

$$b_2 = \ln(O_{\text{Food, Yes}}) - \ln(O_{\text{Food, No}})$$

$$= \ln(28) - \ln(10)$$

$$= 3.332 - 2.303$$

$$= 1.029$$

CHAPTER 18 CATEGORICAL DATA

The important thing is that b_2 is the difference between the log of the observed frequency for cats that received food and danced, and the log of the observed frequency for cats that received food and didn't dance. Put another way, within the group of cats that received food as a reward it represents the difference between cats that didn't dance and cats that did.

Finally, we can look at cats that had affection and danced. In this case, the value of the affection variable is both 1 and the interaction (which is the value of training and the value of dance) is also 1. We can also replace b_0 , b_1 and b_2 with what we now represent. The outcome is the log of the observed frequency for cats that received food and danced (this expected value is 48 – see Table 18.1). Therefore, the equation becomes (I've used the shorthand of A for affection, F for food, Y for yes, and N for no):

$$\ln(O_{A,Y}) = b_0 + (b_1 \times 1) + (b_2 \times 1) + (b_3 \times 1)$$

$$\ln(O_{A,Y}) = b_0 + b_1 + b_2 + b_3$$

$$\ln(O_{A,Y}) = \ln(O_{F,N}) + (\ln(O_{A,N}) - \ln(O_{F,N})) + (\ln(O_{F,Y}) - \ln(O_{F,N})) + b_3$$

$$\ln(O_{A,Y}) = \ln(O_{A,N}) + \ln(O_{F,Y}) - \ln(O_{F,N}) + b_3$$

$$b_3 = \ln(O_{A,Y}) - \ln(O_{F,Y}) + \ln(O_{F,N}) - \ln(O_{A,N})$$

$$= \ln(48) - \ln(28) + \ln(10) - \ln(114)$$

$$= -1.895$$

So, b_3 in the model really compares the difference between affection and food for cats that didn't dance to the difference between food and affection when the cats did dance. Put another way, it compares the effect of training when cats didn't dance to the effect of training when they did dance.

The final model is therefore:

$$\ln(O_{ij}) = 2.303 + 2.433\text{Training}_i + 1.029\text{Dance}_j - 1.895\text{Interaction}_{ij} + \ln(\varepsilon_{ij})$$

The important thing to note here is that everything is exactly the same as in the ANOVA, except that we dealt with log-transformed values (compare this to Section 13.2.2 to see just how similar everything is). In case you still don't believe this works as a general linear model, I've prepared a file called **Cat Regression** which contains the two variables **Dance** (0 = no, 1 = yes) and **Training** (0 = food, 1 = affection) and the interaction (**Interaction**). There is also a variable called **Observed** which contains the observed frequencies in Table 18.1 for each combination of **Training** and **Dance**. Finally, there is a variable called **LnObserved**, which is the natural log of these observed frequencies (remember that throughout this section we've de-

log observed values).



SELF-TEST Run a multiple regression analysis using **Cat Regression.sav** with **LnObserved** as the outcome, and **Training, Dance** and **Interaction** as your three predictors.

OUTPUT 18.1

Coefficients ^a				
		Unstandardized Coefficients		Standardized Coefficients
		B	Std. Error	Beta
1	(Constant)	2.303	.000	
	Type of Training	2.434	.000	1.385
	Did they dance?	1.030	.000	.725
	Interaction	-1.895	.000	-1.174

a. Dependent Variable: LN (Observed Frequencies)

Output 18.1 shows the resulting coefficients table from this regression. One thing to note is that the constant, b_0 , is 2.303 as calculated above, the coefficient for training, b_1 , is 2.434 and for dance, b_2 , is 1.030, both of which are the same as what was calculated above. Also the coefficient for the interaction term is -1.895, which is also the same as what was calculated above. There is one interesting point, though: all of the standard errors are .000, which means there is no error at all in this model (which is also why there are no residuals). This is because the various combinations of coding variables completely predict the observed frequencies. This is known as a **saturated model** – I will return to this point in a moment. For the time being, I hope this convinces you that chi-square can be conceptualized as a linear model.

OK, this is all very well, but the heading of this section did not show you how the chi-square test can be conceptualized as a linear model. The chi-square test looks at whether two variables are independent. In this case, we are interested in the combined effect of the two variables, only their unique effects. We can conceptualize chi-square in much the same way as the saturated model, but we don't include the interaction term. If we remove the interaction term from the model, we get:

$$\ln(\text{model}_{ij}) = b_0 + b_1 \text{Training}_i + b_2 \text{Dance}_j$$

With this new model, we cannot predict the observed values like we could with the saturated model because we've lost some information (namely, the interaction term). The outcome from the model changes, and therefore the beta values change. This is why the chi-square test is based on 'expected frequencies'. Therefore, if we conceptualize the chi-square test as a linear model, our outcomes will be these expected frequencies. Looking back to the beginning of this chapter you'll see we already have expected frequencies based on this model. We can recalculate the beta values based on this model.

$$\ln(E_{ij}) = b_0 + b_1 \text{Training}_i + b_2 \text{Dance}_j$$

For cats that had food reward and didn't dance, the training and dance variables are 0 and so the equation reduces to:

$$\ln(E_{\text{Food, No}}) = b_0 + (b_1 \times 0) + (b_2 \times 0)$$

$$\ln(E_{\text{Food, No}}) = b_0$$

$$b_0 = \ln(23.56)$$

$$= 3.16$$

CHAPTER 18 CATEGORICAL DATA

Therefore, b_0 represents the log of the expected value when all of the categories

When we look at cats that had affection as a reward and didn't dance, the training variable is 1 and the dance variable is still 0. Also, our outcome now changes to be the expected value for cats that received affection and didn't dance:

$$\begin{aligned}\ln(E_{\text{Affection, No}}) &= b_0 + (b_1 \times 1) + (b_2 \times 0) \\ \ln(E_{\text{Affection, No}}) &= b_0 + b_1 \\ b_1 &= \ln(E_{\text{Affection, No}}) - b_0 \\ &= \ln(E_{\text{Affection, No}}) - \ln(E_{\text{Food, No}}) \\ &= \ln(100.44) - \ln(23.56) \\ &= 1.45\end{aligned}$$

The important thing is that b_1 is the difference between the log of the expected frequencies for cats that received affection and didn't dance and the log of the expected frequencies for cats that received food and didn't dance. In fact, the value is the same as the column effect, that is, the difference between the total number of cats getting affection and the total number of cats getting food: $\ln(162) - \ln(38) = 1.45$. Put simply, it represents the effect of the type of training.

When we look at cats that had food as a reward and danced, the training variable is 0 and the dance variable is 1. Our outcome now changes to be the expected frequency for cats that received food and danced:

$$\begin{aligned}\ln(E_{\text{Food, Yes}}) &= b_0 + (b_1 \times 0) + (b_2 \times 1) \\ \ln(E_{\text{Food, Yes}}) &= b_0 + b_2 \\ b_2 &= \ln(E_{\text{Food, Yes}}) - b_0 \\ &= \ln(E_{\text{Food, Yes}}) - \ln(E_{\text{Food, No}}) \\ &= \ln(14.44) - \ln(23.56) \\ &= -0.49\end{aligned}$$

Therefore, b_2 is the difference between the log of the expected frequencies for cats that received food and did or didn't dance. In fact, the value is the same as the row effect, that is the difference between the total number of cats that did and didn't dance: $\ln(124) - \ln(38) = -0.49$. In simpler terms, it is the main effect of whether or not the cat danced.

We can double-check all of this by looking at the final cell (cats that had affection and danced):

$$\ln(E_{\text{Affection, Yes}}) = b_0 + (b_1 \times 1) + (b_2 \times 1)$$

$$\ln(E_{\text{Affection, Yes}}) = b_0 + b_1 + b_2$$

$$\ln(61.56) = 3.16 + 1.45 - 0.49$$

$$4.12 = 4.12$$

The final chi-square model is therefore:

$$\begin{aligned}\ln(O_i) &= \ln(\text{model}) + \ln(\varepsilon_i) \\ &= 3.16 + 1.45\text{Training} - 0.49\text{Dance} + \ln(\varepsilon_i)\end{aligned}$$

We can rearrange this equation to get some residuals (the error term)

$$\ln(\varepsilon_i) = \ln(O_i) - \ln(\text{model})$$

In this case, the model is merely the expected frequencies that were used in the chi-square test, so the residuals are the differences between the observed and expected frequencies.



SELF-TEST To show that this all actually works, we'll run a multiple regression analysis using **Cat Regress**. This time the outcome is the log of expected frequencies (**LnExpected**) and **Training** and **Dance** are the predictors (the interaction is not included).



EVERYBODY

This section demonstrates how chi-square can work as a linear model, and ANOVA, in which the beta values tell us something about the differences in frequencies across categories of our two variables. If nothing else, leave this section aware that chi-square (and analysis of categorical data) can be expressed as a linear model (although we have to use log values for the frequencies of a variable using dummy variables, just as we did with regression, and the resulting beta values can be calculated in exactly the same way as in ANOVA. In ANOVA, these beta values represented differences between a particular category compared against a baseline category. With categorical data, the beta values represent the same thing, the only difference being that regression means we're dealing with expected values. Grasping this idea (that ANOVAs and categorical data analysis are basically the same) will be helpful in the next section.

18.3.6.2. Loglinear analysis ③

In the previous section, after nearly reducing my brain to even more of a mess than it already is trying to explain how categorical data analysis works, I ran the data through an ordinary regression using SPSS. I was talking complete gibberish. At the time I rather glibly said 'oh, by the way, this is a loglinear model'.

in the model, that's odd isn't it?' and sort of passed this off by telling rated' model and not to worry too much about it because I'd explained as I'd worked out what the hell was going on. That seemed like a good time, but unfortunately I now have to explain what I was going

To begin with, I hope you're now happy with the idea that can be expressed in the form of a linear model provided that we use log variables (this is why the technique we're discussing is called *loglinear* analysis). If you already know about ANOVA and linear models generally, you're now tucked up in bed with the idea that we can extend any linear model

CHAPTER 18 CATEGORICAL DATA

of predictors and any resulting interaction terms between predictors. If we cast a simple two-variable categorical analysis in terms of a linear model, then it may amaze you to discover that if we have more than two variables this is no problem: we can extend the model by adding any new variables and the resulting interaction terms, giving each one a parameter (b). This is all you really need to know. So, just as in multiple regression and ANOVA, if we think of things in terms of a linear model, then it becomes very easy to understand how the model expands to incorporate new variables. So, for example, if we have three predictors (A, B and C) in ANOVA we end up with two-way interactions (AB, AC, BC) and one three-way interaction (ABC). The resulting linear model is:

$$\text{outcome}_{ijk} = (b_0 + b_1A_i + b_2B_j + b_3C_k + b_4AB_{ij} + b_5AC_{ik} + b_6BC_{jk} + b_7ABC_{ijk}) +$$

In exactly the same way, if we have three variables in a categorical data analysis and an identical model, but with an outcome in terms of logs:

$$\ln(O_{ijk}) = (b_0 + b_1A_i + b_2B_j + b_3C_k + b_4AB_{ij} + b_5AC_{ik} + b_6BC_{jk} + b_7ABC_{ijk}) +$$

Obviously the calculation of beta values and expected values from the model becomes considerably more cumbersome and confusing, but that's why we invented computers. Now we don't have to worry about it. Loglinear analysis works on these principles. In the way we've seen in the two-variable case, when our data are categorical and we include all the available terms (main effects and interactions) we get no error: our predictors completely predict our outcome (the expected values). So, if we start with the most complex model possible, we will get no error. The job of loglinear analysis is to try to fit a simpler model to the data without any substantial loss of predictive power. Therefore, loglinear analysis typically works on a principle of backward elimination (yes, the same kind of backward elimination as in multiple regression – see Section 8.5.1.3). So we begin with the most complex model, and then remove a predictor from the model and using this new model we fit our data (calculate expected frequencies, just like the chi-square test) and then see if the model fits the data (i.e., are the expected frequencies close to the observed frequencies?). If the fit of the new model is not very different from the more complex model, we abandon the complex model in favour of the new one. Put another way, we assume the term we removed was not having a significant impact on the ability of our model to predict the observed data.

However, we don't just remove terms randomly, we do it hierarchically. We start with the saturated model and then remove the highest-order interaction, and then see if the model fits the data. If removing the interaction term has no effect on the model fit, we assume it is not having much of an effect; therefore, we get rid of it and move on to the next highest-order interaction. If removing these interactions has no effect on the model fit, we move on to any main effects until we find an effect that does affect the fit of the model. The term removed is then considered to be of no significant impact on the model.

Therefore, this is a means of removing terms from the beginning of the equation and looking

To put this in more concrete terms, at the beginning of the section on logit we asked you to imagine we'd extended our training and line-dancing example to a sample of dogs. So, we now have three variables: **Animal** (dog or cat), **Train** (affection) and **Dance** (did they dance or not?). Just as in ANOVA this results in effects:

- **Animal**
- **Training**
- **Dance**

three interactions involving two variables:

- **Animal × Training**
- **Animal × Dance**
- **Training × Dance**

and one interaction involving all three variables:

- **Animal × Training × Dance**

When I talk about backward elimination I mean that loglinear a ing all of these effects; we then take the highest-order interaction way interaction of **Animal × Training × Dance**) and remove it. We without this interaction, and from the model calculate expected the computer) then compare these expected frequencies (or mo observed frequencies using the standard equation for the likelih Section 18.3.3). If the new model significantly changes the likelih removing this interaction term has a significant effect on the fit of t that this effect is statistically important. If this is the case then we that we have a significant three-way interaction. We won't test an all lower-order effects are consumed within higher-order effects. If, three-way interaction doesn't significantly affect the fit of the mod lower-order interactions. Therefore, we look at the **animal × traini training × dance** interactions in turn and construct models in whi present. For each model the computer again calculates expected val to the observed data using a likelihood ratio statistic.² Again, if a results in a significant change in the likelihood ratio then the term is move on to look at any main effects involved in that interaction (so interaction is significant the computer won't look at the main effi ing). However, if the likelihood ratio is unchanged then the analysis interaction term and moves on to look at main effects.

I mentioned that the likelihood ratio statistic (see Section 18.3.3 model. It should be clear how this equation can be adapted to fit a values are the same throughout, and the model frequencies are simply t from the model being tested. For the saturated model, this statistic will observed and model frequencies are the same so the ratio of observed t be 1, and $\ln(0) = 0$), but as we've seen, in other cases it will provide a model fits the observed frequencies. To test whether a new model has ratio, all we need do is to take the likelihood ratio for a model and su hood ratio for the previous model (provided the models are hierarchic

$$L\chi^2_{\text{Change}} = L\chi^2_{\text{Current Model}} - L\chi^2_{\text{Previous Model}}$$

I've tried in this section to give you a flavour of how loglinear a

getting too much into the nitty-gritty of the calculations. I've tried to conceptualize a chi-square analysis as a linear model and then repeatedly told you about ANOVA to hope that you can extrapolate and understand roughly what's going on. The curious among you might

² It's worth mentioning that for every model, the computation of expected values is more complex, the computation gets increasingly tedious and incomprehensible (and you don't need to know the calculations to get a feel for what is going on).

CHAPTER 18 CATEGORICAL DATA

how everything is calculated and to these people I have two things to say: ‘I know a really good place where you can buy a straitjacket’. If you’re then Tabachnick and Fidell (2007) has a wonderfully detailed and lucid chapter on this subject, which puts this feeble attempt to shame.

18.4. Assumptions when analysing categorical data ①

It should be obvious that the chi-square test does not rely on the assumption in Chapter 5 (for example, categorical data cannot have a normal sampling distribution because they aren’t continuous). However, the chi-square test still has two assumptions relating to (1) independence and (2) expected frequencies.

18.4.1. Independence ①

Pretty much all of the tests we have encountered in this book have made an assumption about the independence of residuals, and the chi-square test is no exception. For the chi-square test to be meaningful it is imperative that each person, item or entity belongs to only one cell of the contingency table. Therefore, you cannot use a chi-square test in a repeated-measures design (e.g., if we had trained some cats with food to see if they would dance and then trained the same cats with affection to see if they would dance and then analyse the resulting data with Pearson’s chi-square test).

18.4.2. Expected frequencies ①

With 2×2 contingency tables (two categorical variables both with two categories) the expected values should be below 5. In larger tables, and when looking at associations between three or more categorical variables (loglinear analysis, covered in Section 18.5) the rule is that all expected counts should be greater than 1 and no more than 5. Howell (2012) gives a nice explanation of how violating this assumption creates problems. If this assumption is broken the result is a reduction in test power – so dramatic, in fact, that it may not be worth bothering with the analysis at all.

In terms of remedies, if you’re looking at associations between only two variables you can consider using Fisher’s exact test (Section 18.3.2). With three or more variables (loglinear analysis) your options are: (1) collapse the data across one of the variables (usually the one you least expect to have an effect); (2) collapse levels of one of the variables; (3) collect more data; or (4) accept the loss of power. If you want to collapse

(5) collect more data, or (4) accept the loss of power. If you want to collapse one of the variables then:

- 1 The highest-order interaction should be non-significant.
- 2 At least one of the lower-order interaction terms involving the variable to be collapsed should be non-significant.

Let's think about our loglinear example in which we're looking at the interaction between training (food vs. affection), whether the animal danced (yes vs. no), and whether the animal was a dog or a cat.