

Exploratory Factor Analysis

Theory and Application

1. INTRODUCTION

Many scientific studies are featured by the fact that “numerous variables are used to characterize objects” (Rietveld & Van Hout 1993: 251). Examples are studies in which questionnaires are used that consist of a lot of questions (variables), and studies in which mental ability is tested via several subtests, like verbal skills tests, logical reasoning ability tests, etcetera (Darlington 2004). Because of these big numbers of variables that are into play, the study can become rather complicated. Besides, it could well be that some of the variables measure different aspects of a same underlying variable.

For situations such as these, (exploratory¹) factor analysis has been invented². Factor analysis attempts to bring intercorrelated variables together under more general, underlying variables. More specifically, the goal of factor analysis is to reduce “the dimensionality of the original space and to give an interpretation to the new space, spanned by a reduced number of new dimensions which are supposed to underlie the old ones” (Rietveld & Van Hout 1993: 254), or to explain the variance in the observed variables in terms of underlying latent factors” (Habing 2003: 2) Thus, factor analysis offers not only the possibility of gaining a clear view of the data, but also the possibility of using the output in subsequent analyses (Field 2000; Rietveld & Van Hout 1993).

In this paper an example will be given of the use of factor analysis. This will be done by carrying out a factor analysis on data from a study in the field of applied linguistics, using SPSS for Windows. For this to be understandable, however, it is necessary to discuss the theory behind factor analysis.

2. THE THEORY BEHIND FACTOR ANALYSIS

As the goal of this paper is to show and explain the use of factor analysis in SPSS, the theoretical aspects of factor analysis will here be discussed from a practical, applied perspective. Moreover, already plenty theoretical treatments of factor analysis exist that offer excellent and comprehensive overviews of this technique (see for instance Field 2000: 470; Rietveld & Van Hout 1993: 293-295). It would therefore be useless to come up with another purely theoretical discussion.

The point of departure in this ‘theoretical’ discussion of factor analysis is a flow diagram from Rietveld & Van Hout (1993: 291) that offers an overview of the steps in factor analysis. Treating the theory in such a stepwise and practical fashion should ultimately result in a systematic and practical background that makes it possible to successfully carry out a factor analysis. Before moving on to this, however, it is probably useful to explain very shortly the general idea of factor analysis.

¹ Next to exploratory factor analysis, confirmatory factor analysis exists. This paper is only about exploratory factor analysis, and will henceforth simply be named *factor analysis*.

² A salient detail is that it was exactly the problem concerned with the multiple tests of mental ability that made the psychologist Charles Spearman invent factor analysis in 1904 (Darlington 2004).

2.1. Factor analysis in a nutshell

The starting point of factor analysis is a correlation matrix, in which the intercorrelations between the studied variables are presented. The dimensionality of this matrix can be reduced by “looking for variables that correlate highly with a group of other variables, but correlate very badly with variables outside of that group” (Field 2000: 424). These variables with high intercorrelations could well measure one underlying variable, which is called a ‘factor’.

The obtained factor creates a new dimension “that can be visualized as classification axes along which measurement variables can be plotted” (Field 2000: 424; see for example figure 2 on page 7 of this paper). This projection of the scores of the original variables on the factor leads to two results: *factor scores* and *factor loadings*. Factor scores are “the scores of a subject on a [...] factor” (Rietveld & Van Hout 1993: 292), while factor loadings are the “correlation of the original variable with a factor” (ibid: 292). The factor scores can then for example be used as new scores in multiple regression analysis, while the factor loadings are especially useful in determining the “substantive importance of a particular variable to a factor” (Field 2000: 425), by squaring this factor loading (it is, after all, a correlation, and the squared correlation of a variable determines the amount of variance accounted for by that particular variable). This is important information in interpreting and naming the factors.

Very generally this is the basic idea of factor analysis. Now it is time for a more thorough discussion of factor analysis.

2.2. A stepwise treatment of factor analysis

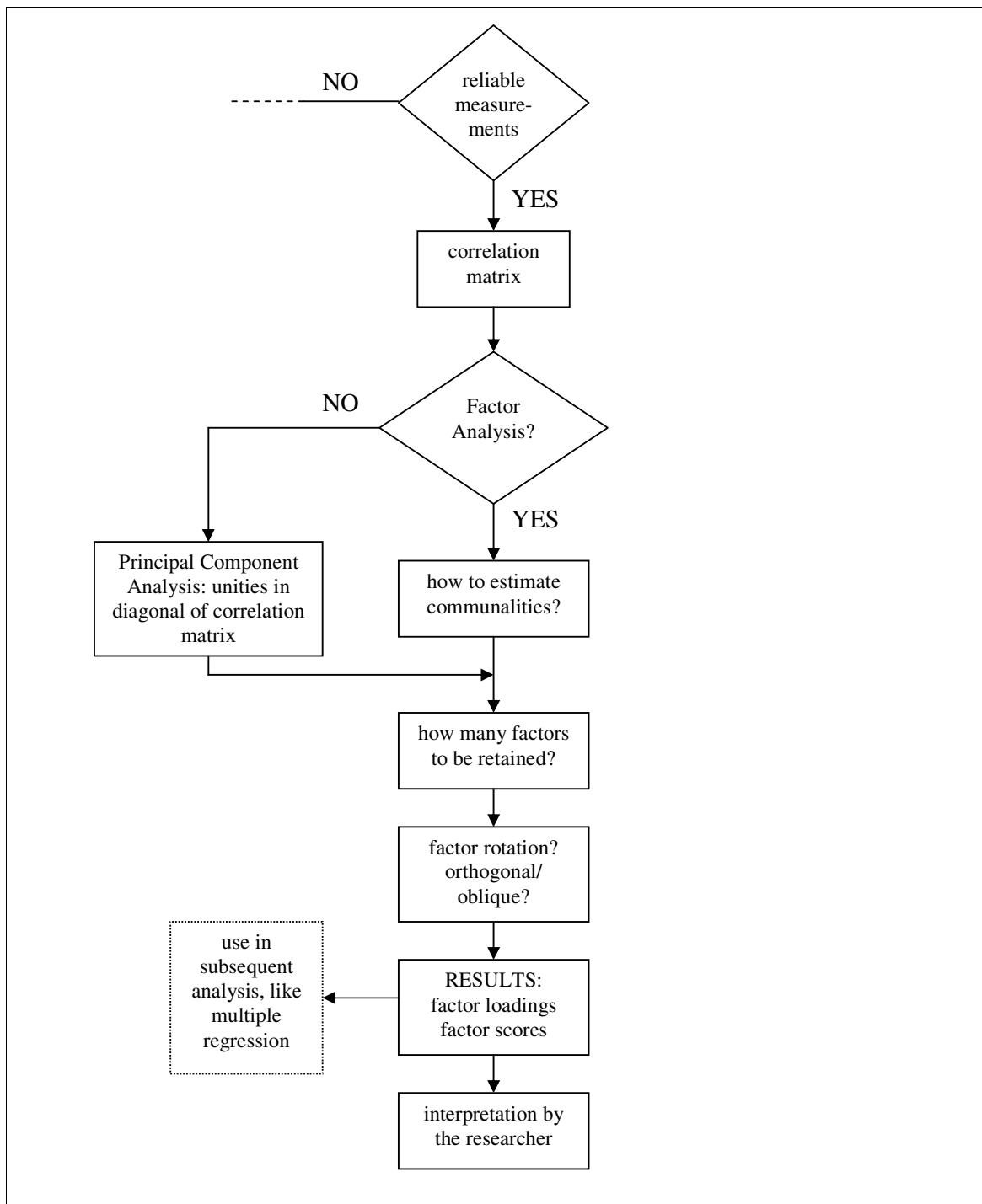
The flow diagram that presents the steps in factor analysis is reproduced in figure 1 on the next page. As can be seen, it consists of seven main steps: reliable measurements, correlation matrix, factor analysis versus principal component analysis, the number of factors to be retained, factor rotation, and use and interpretation of the results. Below, these steps will be discussed one at a time.

2.2.1. Measurements

Since factor analysis departs from a correlation matrix, the used variables should first of all be measured at (at least) an interval level. Secondly, the variables should roughly be normally distributed; this makes it possible to “generalize the results of your analysis beyond the sample collected” (Field 2000: 444). Thirdly, the sample size should be taken into consideration, as correlations are not resistant (Moore & McCabe 2002: 103), and can hence seriously influence the reliability of the factor analysis (Field 2000: 443; Habing 2003).

According to Field (2000: 443) “much has been written about the necessary sample size for factor analysis resulting in many ‘rules-of-thumb’”. Field himself, for example, states that a researcher should have “at least 10-15 subjects per variable” (p. 443). Habing (2003), however, states that “you should have at least 50 observations and at least 5 times as many observations as variables” (p. 3). Fortunately, Monte Carlo studies have resulted in more specific statements concerning sample size (Field 2000: 443; Habing 2003). The general conclusion of these studies was that “the most important factors in determining reliable factor solutions was the absolute sample size and the absolute magnitude of factor loadings” (Field 2000: 443): the more frequent and higher the loadings are on a factor, the smaller the sample can be.

Field (2000) also reports on a more recent study which concludes that “as communalities become lower the importance of sample size increases” (p. 43). This conclusion is adjacent to the conclusion above, as communalities can be seen as a continuation of factor loadings: the communality of a variable is the sum of the loadings of this variable on all extracted factors (Rietveld & Van Hout 1993: 264). As such, the



• Figure 1: overview of the steps in a factor analysis. From: Rietveld & Van Hout (1993: 291).

communality of a variable represents the proportion of the variance in that variable that can be accounted for by all ('common') extracted factors. Thus if the communality of a variable is high, the extracted factors account for a big proportion of the variable's variance. This thus means that this particular variable is reflected well via the extracted factors, and hence that the factor analysis is reliable. When the communalities are not very high though, the sample size has to compensate for this.

In SPSS a convenient option is offered to check whether the sample is big enough: the Kaiser-Meyer-Olkin measure of sampling adequacy (KMO-test). The sample is adequate if the value of KMO is greater than 0.5. Furthermore, SPSS can calculate an *anti-image matrix* of covariances and correlations. All elements on the diagonal of this matrix should be greater than 0.5 if the sample is adequate (Field 2000: 446).

2.2.2. Correlation matrix

When the data are appropriate, it is possible to create a correlation matrix by calculating the correlations between each pair of variables. The following (hypothetical) matrix offers an example of this:

- *Table 1: a hypothetical correlation matrix.*

$$\begin{pmatrix} 1.00 & & & & & \\ \mathbf{0.77} & 1.00 & & & & \\ \mathbf{0.66} & \mathbf{0.87} & 1.00 & & & \\ 0.09 & 0.04 & 0.11 & 1.00 & & \\ 0.12 & 0.06 & 0.10 & \mathbf{0.51} & 1.00 & \\ 0.08 & 0.14 & 0.08 & \mathbf{0.61} & \mathbf{0.49} & 1.00 \end{pmatrix}$$

In this matrix two clusters of variables with high intercorrelations are represented. As has already been said before, these clusters of variables could well be “manifestations of the same underlying variable” (Rietveld & Van Hout 1993: 255). The data of this matrix could then be reduced down into these two underlying variables or factors.

With respect to the correlation matrix, two things are important: the variables have to be intercorrelated, but they should not correlate too highly (extreme multicollinearity and singularity) as this would cause difficulties in determining the unique contribution of the variables to a factor (Field 2000: 444). In SPSS the intercorrelation can be checked by using Bartlett’s test of sphericity, which “tests the null hypothesis that the original correlation matrix is an identity matrix” (Field 2000: 457). This test has to be significant: when the correlation matrix is an identity matrix, there would be no correlations between the variables. Multicollinearity, then, can be detected via the determinant of the correlation matrix, which can also be calculated in SPSS: if the determinant is greater than 0.00001, then there is no multicollinearity (Field 2000: 445).

2.2.3. Factor analysis versus principal component analysis

After having obtained the correlation matrix, it is time to decide which type of analysis to use: factor analysis or principal component analysis³. The main difference between these types of analysis lies in the way the communalities are used. In principal component analysis it is assumed that the communalities are initially 1. In other words, principal component analysis assumes that the total variance of the variables can be accounted for by means of its components (or factors), and hence that there is no error variance. On the other hand, factor analysis does assume error variance⁴. This is reflected in the fact that in factor analysis the

³ There are more types of factor analysis than these two, but these are the most used (Field 2000; Rietveld & Van Hout 1993).

⁴ In this respect, factor analysis is more correct, as it is normal to split variance up into common variance, unique variance and error variance (see Field 2000: 432; Rietveld & Van Hout 1993: 267).

communalities have to be estimated, which makes factor analysis more complicated than principal component analysis, but also more conservative.

2.2.3.1. Technical aspects of principal component analysis

In order to understand the technical aspects of principal component analysis it is necessary to be familiar with the basic notions of matrix algebra. I assume this familiarity to be present with the readers. If not, chapter six of Rietveld & Van Hout (1993) provides a sufficient introduction.

The extraction of principal components or factors in principal component analysis takes place by calculating the eigenvalues of the matrix. As Rietveld & Van Hout (1993: 259) state, “the number of positive eigenvalues determines the number of dimensions needed to represent a set of scores without any loss of information”. Hence, the number of positive eigenvalues determines the number of factors/components to be extracted. The construction of the factor itself is then calculated via a transformation matrix that is determined by the eigenvectors of the eigenvalues (see Rietveld & Van Hout 1993: 261). After constructing the factors it is possible to determine the factor loadings simply by calculating the correlations between the original variables and the newly obtained factors or components.

Rietveld & Van Hout (1993: 258) furthermore name two criteria for dimensionality reduction in principal component analysis:

1. The new variables (principal components) should be chosen in such a way that the first component accounts for the maximum part of the variance, the second component the maximum part of the remaining variance, and so on.
2. The scores on the new variables (components) are not correlated.

2.2.3.2. Technical aspects of factor analysis

In factor analysis the different assumption with regard to the communalities is reflected in a different correlation matrix as compared to the one used in principal component analysis. Since in principal component analysis all communalities are initially 1, the diagonal of the correlation matrix only contains unities. In factor analysis, the initial communalities are not assumed to be 1; they are estimated (most frequently) by taking the squared multiple correlations of the variables with other variables (Rietveld & Van Hout 1993: 269). These estimated communalities are then represented on the diagonal of the correlation matrix, from which the eigenvalues will be determined and the factors will be extracted. After extraction of the factors new communalities can be calculated, which will be represented in a reproduced correlation matrix.

The difference between factor analysis and principal component analysis is very important in interpreting the factor loadings: by squaring the factor loading of a variable the amount of variance accounted for by that variable is obtained. However, in factor analysis it is already initially assumed that the variables do not account for 100% of the variance. Thus, as Rietveld & Van Hout (1993) state, “although the loading patterns of the factors extracted by the two methods do not differ substantially, their respective amounts of explained variance do!” (p. 372).

2.2.3.3. Which type of analysis to choose?

According to Field (2000: 434) strong feelings exist concerning the choice between factor analysis and principal component analysis. Theoretically, factor analysis is more correct, but also more complicated. Practically, however, “the solutions generated from principal component analysis differ little from those derived from factor analysis techniques” (Field 2000: 434). In Rietveld & Van Hout (1993: 268) this is further specified: “the difference

between factor analysis and principal component analysis decreased when the number of variables and the magnitudes of the factor loadings increased”.

The choice between factor analysis thus depends on the number of variables and the magnitude of the factor loadings. After having made this choice, the question arises how many factors there are to be retained.

2.2.4. Number of factors to be retained

As can be inferred from the last section, the number of factors to be retained is similar to the number of positive eigenvalues of the correlation matrix. This may, however, not always lead to the right solutions, as it is possible to obtain eigenvalue that are positive but very close to zero. Therefore, some rules of thumb have been suggested for determining how many factors should be retained (see Field 2000: 436-437; Rietveld & Van Hout 1993: 273-274):

1. Retain only those factors with an eigenvalue larger than 1 (Guttman-Kaiser rule);
2. Keep the factors which, in total, account for about 70-80% of the variance;
3. Make a scree-plot⁵; keep all factors before the breaking point or elbow.

It is furthermore always important to check the communalities after factor extraction. If the communalities are low, the extracted factors account for only a little part of the variance, and more factors might be retained in order to provide a better account of the variance. As has already been discussed in section 2.2.1., the sample size also comes into play when considering these communalities.

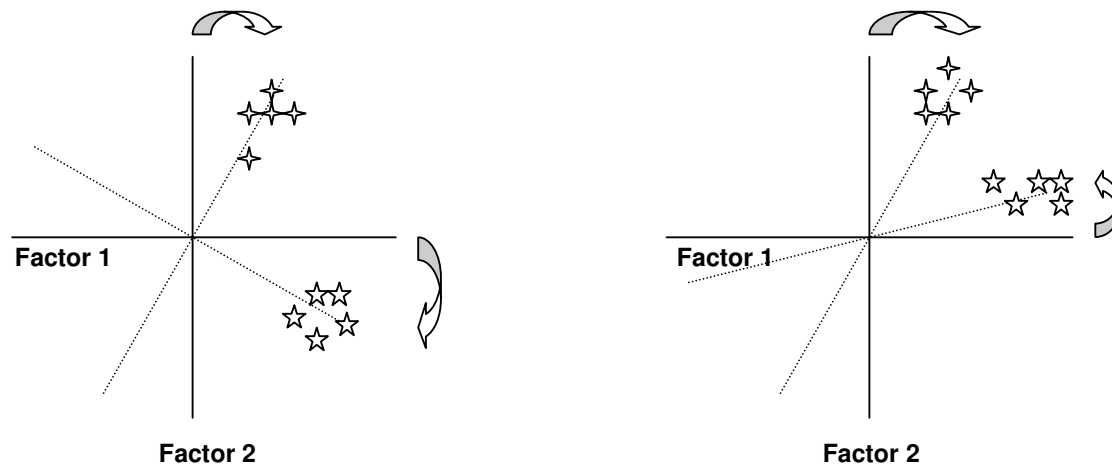
2.2.5. Factor rotation

After factor extraction it might be difficult to interpret and name the factors / components on the basis of their factor loadings. Remember that the criterion of principal component analysis that the first factor accounts for the maximum part of the variance; this will often ensure that “most variables have high loadings on the most important factor, and small loadings on all other factors” (Field 2000: 438). Thus, interpretation of the factors can be very difficult.

A solution for this difficulty is factor rotation. Factor rotation alters the pattern of the factor loadings, and hence can improve interpretation. Rotation can best be explained by imagining factors as axes in a graph, on which the original variables load. By rotating these axes, then, it is possible to make clusters of variables load optimally, as is shown in figure 2 on the next page.

Figure 2 also represents the two types of rotation: orthogonal and oblique rotation. In orthogonal rotation there is no correlation between the extracted factors, while in oblique rotation there is. It is not always easy to decide which type of rotation to take; as Field (2000: 439) states, “the choice of rotation depends on whether there is a good theoretical reason to suppose that the factors should be related or independent, and also how the variables cluster on the factors before rotation”. A fairly straightforward way to decide which rotation to take is to carry out the analysis using both types of rotation; “if the oblique rotation demonstrates a negligible correlation between the extracted factors then it is reasonable to use the orthogonally rotated solution” (Field 2000: 439).

⁵ See Field (2000: 436): A scree plot is “a graph of each eigenvalue (Y-axis) against the factor with which it is associated (X-axis)”.



- *Figure 2: graphical representation of factor rotation. The left graph represents orthogonal rotation and the right one represents oblique rotation. The stars represent the loadings of the original variables on the factors. (Source for this figure: Field 2000: 439).*

There are several methods to carry out rotations. SPSS offers five: varimax, quartimax, equamax, direct oblimin and promax. The first three options are orthogonal rotation; the last two oblique. It depends on the situation, but mostly varimax is used in orthogonal rotation and direct oblimin in oblique rotation. Orthogonal rotation results in a rotated component / factor matrix that presents the ‘post-rotation’ loadings of the original variables on the extracted factors, and a transformation matrix that gives information about the angle of rotation. In oblique rotation the results are a pattern matrix, structure matrix, and a component correlation matrix. The pattern matrix presents the ‘pattern loadings’ (“regression coefficients of the variable on each of the factors”; Rietveld & Van Hout 1993: 281) while the structure matrix presents ‘structure loadings’ (“correlations between the variables and the factors”; *ibid.*); most of the time the pattern matrix is used to interpret the factors. The component correlation matrix presents the correlation between the extracted factors / components, and is thus important for choosing between orthogonal and oblique rotation.

2.2.6. Results: factor loadings and factor scores

In the last paragraph it was already reported that the factor loadings are represented in the rotated component matrix. As may be known by now, these factor loadings are important for the interpretation of the factors, especially the high ones. One can wonder, however, how high a loading has to be in order to determine the interpretation of the factor in a significant way. This is dependent of the sample size (Field 2000: 440): the bigger the sample the smaller the loadings can be to be significant. Stevens (1992) made a critical values table to determine this significance (see Field 2000: 440). Field (2000: 441) states, on the other hand, that “the significance of a loading gives little indication of the substantive importance of a variable to a factor”. For this to determine, the loadings have to be squared. Stevens (1992: in Field 2000: 441) then “recommends interpreting only factor loadings with an absolute value greater than 0.4 (which explain around 16% of variance)”. This is only possible in principal component analysis, though. In factor analysis the amount of explained variance is calculated in a different way (see the last paragraph of section 2.2.3.2.). Thus, Stevens recommendation should be approached with care!

The second result of factor analysis is the factor scores. These factor scores can be useful in several ways. Field (2000: 431) and Rietveld & Van Hout (1993: 289-290) name the following:

1. If one wants to find out “whether groups or clusters of subjects can be distinguished that behave similarly in scoring on a test battery, [and] the latent, underlying variables are considered to be more fundamental than the original variables, the clustering of factor scores in the factor space can provide useful clues to that end” (Rietveld & Van Hout: 289)
2. The factor scores can serve as a solution to multicollinearity problems in multiple regression. After all, the factor scores are uncorrelated (in the case of orthogonal rotation);
3. The factor scores can also be useful in big experiments, containing several measures using the same subjects. If it is already known in advance that “a number of dependent variables used in the experiment in fact constitute similar measures of the same underlying variable, it may be a good idea to use the scores on the different factors, instead of using the scores on the original variables” (Rietveld & Van Hout: 290).

In SPSS the factor scores for each subject can be saved as variables in the data editor. Doing this via the *Anderson-Rubin* method (option in SPSS) ensures that the factor scores are uncorrelated, and hence usable in a multiple regression analysis. If it does not matter whether factor scores are correlated or not, the *Regression* method can be used. The correlation between factor scores can also be represented in a factor score covariance matrix, which is displayed in the SPSS output (next to a factor score coefficient matrix, which in itself is not particularly useful (Field 2000: 467)).

3. FACTOR ANALYSIS IN PRACTICE: AN EXAMPLE

Now that it is clear what the principles of and steps in factor analysis are, it is time to bring it into practice. As has already been said, this will be done by using data from an applied linguistics study. In this case, this is the doctoral research of Sible Andringa (see Andringa 2004).

3.1. An introduction to Andringa's study

Andringa's research is about “the effect of explicit and implicit form-focused instruction on L2 proficiency” (Andringa 2004: 1). Form-focused instruction is a much researched issue in applied linguistics, and refers to “any pedagogical effort which is used to draw the learner's attention to language form either implicitly or explicitly” (Spada 1997: 226). Implicit instruction is mainly meaning-based and is also named *focus on form*. Explicit instruction, on the other hand, is more directed to “discrete-point grammatical presentation and practice” (Spada 1997: 226) and is also called *focus on forms*.

In Andringa's study the question is raised whether “the kind of linguistic knowledge that is developed with form-focused instruction [can] transfer to spontaneous situations of use?” (Andringa 2004: 4). In research on form-focused instruction it was concluded that this transfer is amongst others dependent on three ‘key-factors’: the type of instruction, the nature of the structure to be learnt, and the individual differences between the learners (Andringa 2004). Andringa therefore focuses specifically on these three aspects.

As such, the main features of the study are as follows (Andringa 2004: 12):

- Computer-instructed language learning of Dutch as a second language with performance tracking;
- 3 conditions: explicit Focus on Forms, implicit Focus on Form, and control (no instruction);
- Instruction in two contrasting grammar rules;
- Design: pretest, treatment, posttest and second posttest
- Dependent variables: proficiency in use of grammar rules and declarative knowledge of grammar rules
- Learner variables (independent):
 - writing proficiency measures
 - general L2 proficiency test
 - aptitude: memory and grammatical sensitivity
 - numerous affective variables, like sex and age

This needs some more specification. First of all I will consider the targeted grammar rules: these rules are *degrees of comparison* and *subordination*. Degrees of comparison (*mooi – mooier – mooist*) are morphosyntactic and considered both meaningful and simple (Andringa jaartal). Subordination, on the other hand, has a syntactic basis: *hij is een aardige jongen → omdat hij een aardige jongen is*. This word-order rule in subordination is considered meaningless but simple (ibid.).

Above, the difference between implicit and explicit form-focused instruction has already been explained. In Andringa (2004: 15) the implicit instruction is operationalized by embedding the target structures in text-comprehension exercises: *why does Coca cola advertise? A- they advertise, because they a lot of cola want to sell; B- ...*). The explicit instruction is done via the presentation of the target structures in grammar exercises: *look at the highlighted words: which sentence is correct? A- Coca cola advertises, because they a lot of cola want to sell; B- ...*).

The dependent variables are measured by a written proficiency test and declarative knowledge test for both target structures. The proficiency test is divided into 19 submeasurements per target structure; the declarative knowledge test measures have only been divided into the two target structures⁶. In addition, the declarative knowledge test is only administered as a posttest in order to ensure that no learning effect will have taken place. Both the proficiency test and the declarative knowledge test are included in the appendix.

Finally, the first three learner variables are represented as covariates. The writing proficiency is measured by examining the grammatical accuracy (errors/clause), fluency (words/clause) and spelling (sp.errors/clause) of the written proficiency test. General L2 proficiency is tested via the so-called C-test (see appendix), Cito ISK-test. Aptitude is measured by means of a memory test and a grammatical sensitivity test (see appendix). In addition, a grammatical judgement test has been administered.

From this then, Andringa (2004: 18) formulated the following hypotheses:

- Declarative knowledge:

1. A – Degrees of comparison:	expl. FonFs	>	impl. FonF
B – Subordination:	expl. FonFs	>	impl. FonF
- Implicit knowledge / proficiency:

2. A – Degrees of comparison:	expl. FonFs	<	impl. FonF
B – Subordination:	expl. FonFs	<	impl. FonF

⁶ An obvious reason for this is that in the proficiency test spontaneous language use is tested, while in the declarative knowledge test no spontaneous language is used (see appendix).

In other words, Andringa expects that the explicit instruction is most effective for declarative knowledge and the implicit instruction for proficiency (i.e. spontaneous language use).

3.2. *The use of factor analysis in this study*

After this introduction of Andringa's study, the question arises how factor analysis can be useful in this study. To give an answer to this question it is necessary to go back to section 2.2.6., in which the uses of factor scores were listed. One of these was the use of factor scores in big experiments in which several measurements are taken using the same subjects: when it is known in advance that a number of these measurements actually are similar measurements of the same underlying variable, it might be useful to use the scores on this underlying variable (factor / component) instead of all the scores of the original variables in subsequent analyses such as (M)ANOVA. Remembering that the tests in the study of Andringa are divided into several submeasurements; it would probably be useful to run a factor analysis on these data to check whether it is possible to make use of factor scores.

Because of the design of the study (pretest / posttest), however, this would mean that quite a lot of analyses have to be carried out, namely for the pretest subordination measures, the pretest degrees of comparison measures, the posttest subordination measures in condition 1, 2, and 3 separately, and the posttest degrees of comparison measures in condition 1,2, and 3 separately. This would be rather complicated, next to the fact that the division in the three conditions would make the sample rather small, and hence the analysis less reliable.

Another interesting use of factor analysis in this study is to determine how the covariates covary with the dependent variables. As it happens, these covariates measure different aspects of the learner's knowledge, like aptitude, general language proficiency, and fluency. By discovering which submeasurements cluster on which covariates to which account, it might be possible to find out which kinds of knowledge are used in which submeasurements. This could be important for the choice of instruction type. In the remaining part of this paper this last use of factor analysis will be brought into practice.

3.3. *The analysis*

Although some aspects of the running of factor analysis in SPSS have already been touched upon above, it might be useful to elaborate a bit more on the SPSS-procedure of factor analysis before the analysis will be carried out.

3.3.1. *Factor analysis in SPSS*

Factor analysis can be run in SPSS by clicking on *Analyze* → *Data reduction* → *Factor*. A dialog box will be displayed in which the variables to be analyzed can be selected. Next to this there are five buttons which more or less are similar to the steps in factor analysis described above: *descriptives*, *extraction*, *rotation*, *scores*, and *options*. The *descriptives* button gives the option to select descriptive statistics and the correlation matrix (with the coefficients, significance levels, determinant, KMO & Bartlett's test of sphericity, inverse, reproduced matrix, and anti-image matrix). Clicking on the *extraction* button offers the possibility to decide both which type of analysis to use (factor or principal components) and the number of factors to extract, and to select the scree-plot. Via the *rotation* button, then, the type of rotation can be selected, while the 'scores' menu gives the possibility to save the scores as variables in the data editor and display the factor score coefficient matrix. Finally, the *options* menu gives the possibility to select the way of excluding missing values and to both sort the coefficient display by size and suppress the factor loadings in the factor / component matrix that are lower than a chosen value (i.e. here it is possible to decide how high the value of a factor loading must be to play a role in the interpretation of the factors).

3.3.2. The analysis using Andringa's data⁷

The variables that are included in this factor analysis are the covariates and the dependent variables of the pretest: 23 proficiency test measurements (subordination measures as well as degrees of comparison) and 11 covariates are included. In total, 54 subjects are included (originally 102 subjects participated, but after listwise case exclusion 54 subjects remained). In the descriptives menu all options have been selected.

The next step is the type of factor extraction. I tried both principal component analysis and factor analysis, but because of data problems (on which I will come later on) only principal component analysis worked (the communalities were, by the way, high enough to opt for principal component analysis). Furthermore, I initially selected to extract the factors with eigenvalues over 1 (this is default in SPSS).

Finally, I selected varimax rotation in the rotation menu, chose to display the factor score coefficient matrix in the scores menu and opted for listwise exclusion, sorting by size and suppression of absolute values less than 0.50 in the options menu. I have chosen for a value of 0.50 because the sample is not very big (see Field 2000: 440).

The result of this factor analysis was rather disappointing: the correlation matrix was 'not positive definite', and the determinant was smaller than 0.00001. A logical conclusion, then, is that the data is not appropriate for a factor analysis because of multicollinearity (see section 2.2.2.), and no sound conclusions can be drawn from the analysis. Nonetheless, SPSS ran the factor analysis, and therefore I will interpret this output anyway, keeping in mind that conclusions can only be very tentative.

The result of the analysis was a rotated component matrix consisting of ten components that account for 81,77% of the variance. The breaking point of the scree-plot, however, was at 5 or 6 factors; extraction of only 6 factors would account for only 63,51% of the variance, but enhances the overview of the rotated component matrix considerably. Therefore, I carried out the factor analysis another time, choosing for the extraction of only 6 factors. Moreover, I also ran the factor analysis using oblique rotation (direct oblimin; firstly with factor extraction via eigenvalues >1 and then via the extraction of 6 factors), as I could not think of a reason why the factors would not be correlated. Ultimately, then, I did four factor analyses (see *output_FA* on the floppy in the appendix).

In my view, the oblique rotation with six extracted factors gives the best possibility to interpret the factor solution (although it is important to keep in mind that this solution accounts for only 68% of the variance). The choice of oblique rotation in favour of orthogonal rotation is because of the fact that the factor scores in oblique rotation correlate rather highly, as the component score covariance matrix makes clear:

Component Score Covariance Matrix

Component	1	2	3	4	5	6
1	1,263	-,551	1,993	-,161	-,450	1,935
2	-,551	1,165	-,270	,101	2,053	-,822
3	1,993	-,270	3,078	-,132	1,587	1,492
4	-,161	,101	-,132	1,273	-,243	1,607
5	-,450	2,053	1,587	-,243	4,021	-,843
6	1,935	-,822	1,492	1,607	-,843	4,063

Extraction Method: Principal Component Analysis.

Rotation Method: Oblimin with Kaiser Normalization.

⁷ The data, next to the complete output of the factor analysis, can be examined on the floppy disk in the appendix (in the envelope).

In oblique rotation, interpretation of the factors mostly takes place by examining the pattern matrix (see section 2.2.5.). Therefore the pattern matrix is represented below:

Pattern Matrix^a

	Component					
	1	2	3	4	5	6
Tot vergr. tr. correct Pretrap	,928					
Tot correct lexicaal Pretrap	,854					
Tot vergrot. trap Pretrap	,837					
Tot correct Pretrap	,828					
Tot correct gespeld Pretrap	,786					
Gram. Judg. trappen	,725					
Cito ISK toets	,608					
Tot approp. contexts Pretrap	,514					
C-toets						
Tot non-use in approp. context Pretrap						
Tot cor os Presub		-,869				
Tot cor os other Presub		-,808				
Gram. Judg. subord		-,687				
Tot types os Presub		-,687				
Written Gram err / clause (correction tr)		,677				
Written Gram err / clause (correction os)		,656				
Tot cor os "als" Presub		-,650				
Tot approp. contexts Presub		-,513				
Tot overtr. tr. Pretrap			,891			
Tot overtr. tr. correct Pretrap			,838			
Written Words/Clause			,574			
Tot incorrect Pretrap						
Written Spell errors/ clause				,919		
Written spell err / tot words				,918		
Aptitude Memory				-,610		
Aptitude Analysis						
Contrasten Pretrap						
Tot inc os Presub					,896	
Tot inc os other					,895	
Tot non-use in approp. context Presub					-,611	
Verdachte -e Pretrap						,743
Avoided Presub						
Tot inc os "als" Presub						
Avoided Pretrap						

Extraction Method: Principal Component Analysis.

Rotation Method: Oblimin with Kaiser Normalization.

a. Rotation converged in 33 iterations.

The first factor consists mainly of degrees of comparison measurements ('trappen'), but the ISK-test of the Cito also loads high on this factor. This could mean that for the ISK-test the same kind of knowledge is necessary as for degrees of comparison assignments. The second factor contains subordination measures and grammatical errors; the subordination measures correlate negatively to the grammatical errors. In the third factor, a fluency measure is clustered with measures of superlative in degrees of comparison. It could be then that the use of superlative mostly occurs when the language learners are relatively fluent. In the fourth factor, spelling errors and memory correlate negatively, which implicates in this case that the learners with better memory make fewer spelling errors. This seems rather plausible. The fifth factor constitutes subordination measures: incorrect usage correlates negatively to non-use. Thus, when the non-use does not occur, incorrect usage occurs. This probably takes place in difficult exercises. Finally, the sixth factor simply constitutes one variable: 'verdachte -e' in degrees of comparison (*mooie instead of mooi or mooier*).

From this my suggestions for the names of the factors are:

1. 'General language proficiency' (because of the loading of the ISK-toets)
2. 'Grammatical proficiency'
3. 'Fluency'
4. '(Language) Memory'
5. 'Behaviour in difficult exercises'
6. 'Suspected -e'

As such, it was possible to examine which kinds of knowledge come into play in which measurements. Furthermore, it can be inferred that subordination and degrees of comparison are indeed different kinds of rules, as they load separately on different factors. For the future, it might be interesting to carry out a similar analysis using other and more target structures. After all, if it is known which kind of knowledge are important for learning grammatical structures, then these kinds of knowledge can be emphasized in teaching a particular target structure. Thus, language teaching can become more effective.

4. CONCLUSION

In this paper an example is given of the use of factor analysis in an applied linguistics study. Unfortunately, the data was not appropriate because of multicollinearity, and hence no sound conclusions can be drawn from this analysis. Nevertheless, a principal component analysis has been carried out with oblique rotation. This resulted into six correlated factors, constituting several aspects of 'knowledge' in language learning. It turned out that the measurements of the two target structures loaded on different factors, which could indicate that different kinds of knowledge are needed for these structures. This can be important for language teaching, as it gives opportunities to improve language teaching effectiveness.

BIBLIOGRAPHY

- Andringa, S. (2004). *The effect of Explicit and Implicit Form-Focused Instruction on L2 Proficiency*. PowerPoint presentation of lecture at the AAAL conference 2004 (not published).
- Darlington, R.B. (2004). *Factor Analysis*. Website:
<http://comp9.psych.cornell.edu/Darlington/factor.htm> (accessed 10 May 2004; publication year not clear, therefore the date of access between brackets).
- Field, A. (2000). *Discovering Statistics using SPSS for Windows*. London – Thousand Oaks – New Delhi: Sage publications.
- Habing, B. (2003). *Exploratory Factor Analysis*. Website:
<http://www.stat.sc.edu/~habing/courses/530EFA.pdf> (accessed 10 May 2004).
- Moore, D.S. & McCabe, G.P. (2001). *Statistiek in de Praktijk. Theorieboek*. Schoonhoven: Academic Services
- Rietveld, T. & Van Hout, R. (1993). *Statistical Techniques for the Study of Language and Language Behaviour*. Berlin – New York: Mouton de Gruyter.
- Spada, N. (1997). Form-Focussed Instruction and SLA: a review of classroom and laboratory research. *Language Teaching* 30 (2): 73-87.
- Stevens, J.P. (1992). *Applied Multivariate Statistics for the Social Sciences* (2nd edition). Hillsdale, NJ: Erlbaum.

APPENDIX:

- **Proficiency test**
- **Declarative knowledge test**
- **C-test**
- **Memory test**
- **Grammatical sensitivity test**
- **Grammatical judgement test**
- **Floppy disk containing the data + output (see envelope)**