

# Chapter 9

## Factor Analysis

### Factor Analysis and Principal Component Analysis

The principal component analysis is used for data reduction. Similar to principal component analysis, factor analysis is a multivariate method used for data reduction. The idea is to represent a set of variables by a smaller number of variables, called factors. Factor analysis represents the relationships among multiple correlated/interconnected variables by identifying a few underlying factors.

Factor analysis differs from analysis of variance, multiple linear regression, and discriminant analysis in that, in all such techniques, one variable is considered dependent and the others independent, whereas factor analysis does not make such a distinction. Factor analysis helps explore the possible interrelationships among multiple variables and assess the underlying causes of these relationships.

In factor analysis, the observed variables are expressed as linear functions of a smaller number of underlying factors, analogous to a multiple linear regression model. The key difference is that the explanatory variables, called factors, are not directly observed but are latent variables that are inferred from the correlations among the observed variables.

The model in factor analysis is based on the observed vector, which is partitioned into an unobserved systematic part and an unobserved error part. The systematic part is taken as a linear combination of a relatively small number of unobserved factor variables. The components of error vectors are considered as uncorrelated or independent.

Factor analysis reveals the underlying factors that explain the correlations among variables. It also helps create a smaller set of uncorrelated variables, simplifying further multivariate analyses such as discriminant or regression analysis. It helps identify a smaller and more meaningful set of variables for further use in multivariate analysis.

The applications of factor analysis are exploratory in nature, so it is also called as exploratory factor analysis (EFA).

The essential purpose is to describe the covariance relationships among many variables in terms of a few underlying, but unobservable random quantities called factors. The motivation behind the factor model is that, suppose variables can be grouped by their correlations, such that all variables within a group are highly correlated among themselves but have relatively small correlations with variables in a different group. It can be inferred that each group of variables represents a single underlying construct (or factor) that accounts for the observed correlations. Factor analysis was originally developed for analysing scores on mental tests. For example, when a battery of tests is given to a group of individuals, it is observed that the score of an individual on a given test is more related to his scores on other tests than to the scores of other individuals on the other tests, i.e., usually the scores for any particular individual are interrelated to some degree. This interrelation is “explained” by condensing a test score of an individual as made up of a part which is peculiar to this particular test (called error) and a part which is a function of more fundamental quantities called factor scores. Since they enter several test scores, it is their effect that connects the various test scores. The idea behind this is that a person who is more intelligent in some respects will do better on many tests than someone who is less intelligent.

Both factor analysis and principal component analysis aim to reduce the dimensionality of the data, but they differ in their underlying approaches. Principal component analysis constructs new variables, called principal components, which are linear combinations of the observed variables. In contrast, factor analysis models the observed variables as linear functions of a smaller number of underlying latent factors.

## The Model

The factor model can be considered in the term of a multiple linear regression model where the observable variable  $X_i$  is to be predicted from the values of unobservable common factors  $f_i$  as follows:

$$\begin{aligned}X_1 &= \mu_1 + \lambda_{11}f_1 + \lambda_{12}f_2 + \dots + \lambda_{1m}f_m + \epsilon_1 \\X_2 &= \mu_2 + \lambda_{21}f_1 + \lambda_{22}f_2 + \dots + \lambda_{2m}f_m + \epsilon_2 \\&\vdots \\X_p &= \mu_p + \lambda_{p1}f_1 + \lambda_{p2}f_2 + \dots + \lambda_{pm}f_m + \epsilon_p\end{aligned}$$

where  $E(X_i) = \mu_i$  is the population mean of variable  $i$  and can be viewed as intercept terms in the multiple regression model. The  $i^{th}$  common factor is denoted by  $f_i$  and generally  $m \ll p$ . These variables can be expressed as

$$\underset{\sim}{X} = \underset{\sim}{\mu} + \underset{\sim}{\Lambda}\underset{\sim}{f} + \underset{\sim}{\epsilon}$$

where  $\underset{\sim}{\mu} = (\mu_1, \mu_2, \dots, \mu_p)'$  is a  $p \times 1$  vector of population mean,  $\underset{\sim}{f} = (f_1, f_2, \dots, f_m)'$  is an  $m \times 1$  vector of common factors. The loading of the  $i^{th}$  variable on the  $j^{th}$  factor is denoted by  $l_{ij}$  and are represented in the following  $p \times m$  matrix

$$\underset{\sim}{\Lambda} = \begin{pmatrix} \lambda_{11} & \lambda_{12} & \dots & \lambda_{1m} \\ \lambda_{21} & \lambda_{22} & \dots & \lambda_{2m} \\ \vdots & \vdots & \ddots & \vdots \\ \lambda_{p1} & \lambda_{p2} & \dots & \lambda_{pm} \end{pmatrix}$$

as the matrix of factor loadings, the errors  $\epsilon_i$  is called the specific factor for the variable  $i$  and  $\underset{\sim}{\epsilon} = (\epsilon_1, \epsilon_2, \dots, \epsilon_p)'$  is a  $p \times 1$  vector of specific factors. Note that  $l_{ij}$ 's are like regression coefficients and each of the response variables  $X$  is predicted as a linear function of the unobserved common factors  $f_1, f_2, \dots, f_m$ , which are  $m$  explanatory variables. Thus  $m$  unobserved factors control the variation in the data.

## Assumptions

We assume  $E(\underline{\epsilon}) = 0$ ,  $E(\underline{\epsilon}\underline{\epsilon}') = \Psi$  is a diagonal matrix and  $\underline{\epsilon}$  is distributed independently of  $\underline{f}$ . The vector  $\underline{f}$  is treated as a random vector and as a vector of parameters that varies from observation to observation.

In terms of mental tests, each component of  $X$  is a score on a test and corresponding component of  $\underline{\mu}$  is the average score of this test in the population. The components of  $\underline{f}$  are the scores of the mental factors, and the linear combinations of these enter into the test scores. The coefficients of these linear combinations are the factor loadings, which are the elements of  $\Lambda$ .

The elements of  $\underline{f}$  are also called as common factors because they are common to several different tests. The components of  $\underline{\epsilon}$  are the part of the test score not explained by the common factors and is considered as error of measurement in the test, along with a specific factor having to do only with this particular test. Since in our model, with one set of observations on each individual, we cannot distinguish between these two components of  $\underline{\epsilon}$ , so we call them as error of measurement.

Based on the nature of  $\underline{f}$ , we can have two types of models.

**Model 1:** We consider  $\underline{f}$  to be nonrandom. The components of  $\underline{f}$  varies from one individual to another. It is more accurate to write

$$X_{\alpha} = \mu + \Lambda f_{\alpha} + \epsilon$$

where the nonrandom factor score may describe the systematic part better, but the likelihood function may not have a maximum, which may pose problems of inference. It is suitable when the specific individuals involved and not just the structure are of interest.

**Model 2:** We consider  $\underline{f}$  to be random, which is more appropriate when different samples consist of different individuals.

When  $\underline{f}$  is taken as random, we assume  $E(\underline{f}) = 0$ .

When  $\underline{f}$  is taken as nonrandom, we assume  $E(\underline{X}) = \underline{\mu} + \Lambda E(\underline{f})$  and  $\underline{\mu}$  can be redefined

to absorb  $\Lambda E(\tilde{f})$  and the model is usually expressed as

$$\tilde{X} - \tilde{\mu} = \Lambda \tilde{f} + \tilde{\epsilon}.$$

Let  $E(\tilde{f}\tilde{f}') = \Phi$  when  $\tilde{f}$  is random and usually,  $\tilde{f}$  and  $\tilde{\epsilon}$  are considered to have normal distributions.

If  $f$  is not random, then  $\tilde{f} = f_\alpha$  for the  $\alpha^{th}$  individual and then we assume

$$\frac{1}{N} \sum_{\alpha=1}^N \tilde{f}_\alpha = 0 \quad \text{and} \quad \frac{1}{N} \sum_{\alpha=1}^N \tilde{f}_\alpha \tilde{f}_\alpha' = \Phi.$$

## Indeterminacy in the Model

Let  $C$  be an  $m \times m$  nonsingular matrix. Then let  $\tilde{f}^* = C^{-1}\tilde{f}$  and  $\Lambda^* = \Lambda C$  and the model

$$\tilde{X} = \tilde{\mu} + \Lambda \tilde{f} + \tilde{\epsilon} \tag{1}$$

can be written as

$$\begin{aligned} \tilde{X} &= \tilde{\mu} + \Lambda C C^{-1} \tilde{f} + \tilde{\epsilon} \\ &= \tilde{\mu} + \Lambda^* \tilde{f}^* + \tilde{\epsilon} \end{aligned} \tag{2}$$

When  $\tilde{f}$  is random, then

$$E(\tilde{f}^* \tilde{f}^{*'}) = C^{-1} E(\tilde{f} \tilde{f}') (C^{-1})' = C^{-1} \Phi (C^{-1})' = \Phi^*.$$

When  $\tilde{f}$  is nonrandom, then

$$\frac{1}{N} \sum_{\alpha=1}^N \tilde{f}_\alpha^* \tilde{f}_\alpha^{*'} = \Phi^*.$$

The models (1) with  $\Lambda$  and  $\tilde{f}$  is equivalent to model (2) with  $\Lambda^*$  and  $\tilde{f}^*$  in the sense that by observing  $\tilde{X}$ , we cannot distinguish between model (1) and model (2).

Some of the indeterminacy can be eliminated by requiring

$$E(\tilde{f} \tilde{f}') = I \quad \text{if } \tilde{f} \text{ is random}$$

or

$$\sum_{\alpha=1}^N \tilde{f}_{\alpha} \tilde{f}'_{\alpha} = NI \quad \text{if } \tilde{f} \text{ is nonrandom.}$$

In this case, the factors are said to be **orthogonal**.

If  $\Phi$  is not diagonal, the factors are said to be **oblique**.

When  $\Phi = I$ , then

$$E(\tilde{f}^* \tilde{f}^{*'}) = C^{-1}(C^{-1})' = I \quad (I = CC').$$

The indeterminacy is equivalent to multiplication by an orthogonal matrix. This is called as **problem of rotation**.

Requiring  $\Phi$  to be diagonal needs components of  $\tilde{f}$  to be independently distributed when  $\tilde{f}$  is normally distributed. This condition appeals to psychologists because, by definition, one idea about common mental factors is that they are independent or uncorrelated quantities.

The assumption that the components of  $\xi$  are uncorrelated is crucial. One point of view is that the errors of observation and the specific factors are, by definition, uncorrelated; i.e., the interrelationships among the test scores are caused by common factors that we want to investigate. Another point of view is that the common factors are supposed to explain as much of the variance of the test scores as possible.

If the factors are random, then the covariance matrix of observed  $X$  is

$$\begin{aligned} \Sigma &= E(\tilde{X} - \mu)(\tilde{X} - \mu)' \\ &= E(\Lambda \tilde{f} + \xi)(\Lambda \tilde{f} + \xi)' \\ &= \Lambda \Phi \Lambda' + \Psi. \end{aligned}$$

In case factors are orthogonal (i.e.,  $E(\tilde{f} \tilde{f}') = I$ ), then

$$\Sigma = \Lambda \Lambda' + \Psi.$$

If  $\tilde{f}$  and  $\xi$  are normal, all the information about the structure comes from  $\Sigma$  and  $E(\tilde{X}) = \mu$ .

## Identification

Consider the following question: Given  $\Sigma (= \Lambda\Phi\Lambda' + \Psi)$  and a number  $m$  of factors, can we find whether there exists a triplet

$$(\Lambda, \Phi > 0, \Psi > 0 \text{ and diagonal})$$

to satisfy  $\Lambda\Phi\Lambda' + \Psi = \Sigma$ , and if such a triplet is found, is it unique?

Since any triplet can be transformed into an equivalent structure  $\Lambda C$ ,  $C^{-1}\Phi C^{-1'}$  and  $\Psi$ , we can put  $m^2$  independent conditions on  $\Lambda$  and  $\Phi$  to rule out this indeterminacy.

The number of parameters in

- $\Lambda$  is  $pm$ .
- $\Phi$  is  $\frac{m(m+1)}{2}$  and
- $\Psi$  is  $p$ .

So total number of parameters is  $pm + \frac{m(m+1)}{2} + p = Q$ , say.

The number of components in the observable  $\Sigma$  and the number of conditions (for uniqueness) is  $\frac{p(p+1)}{2} + m^2 = R$ , say.

The excess of observed quantities and conditions over the number of parameters is  $R - Q = \frac{1}{2}[(p - m)^2 - p - m]$ .

- If this excess is positive, one can expect a problem of existence, but can anticipate uniqueness if a set of parameters does exist.
- If this excess is negative, we can expect existence but possibly not uniqueness.
- If this excess is zero, we can hope for both existence and uniqueness (or at least a finite number of solutions).

The question of the existence of a solution is whether there exists a diagonal matrix  $\Psi$  with nonnegative diagonal entries such that  $(\Sigma - \Psi)$  is positive semidefinite matrix of rank  $m$ .

The model is said to be **identified** if a solution exists and is unique.

Some  $m^2$  conditions have to be put on  $\Lambda$  and  $\Phi$  to eliminate a transformation  $\Lambda^* = \Lambda C$  and  $\Phi^* = C^{-1}\Phi C^{-1'}$ . The condition  $\Phi = I$  forces a transformation  $C$  to be orthogonal and there are  $\frac{m(m+1)}{2}$  component equations in  $\Phi = I$ .

Sometimes, the restrictions  $\Gamma = \Lambda'\Psi^{-1}\Lambda$  is diagonal is added. If the diagonal elements of  $\Gamma$  are ordered and different ( $\gamma_{11} > \gamma_{22} > \dots > \gamma_{mm}$ ),  $\Lambda$  is uniquely determined.

## Units of Measurements

In many cases, the units of measurement of each component in  $\underline{X}$  are different. If the units of measurement are to be changed, it means multiplying each component of  $\underline{X}$  by a constant, and these constants may not necessarily be equal. When a given test score is multiplied by a constant, the factor loadings for the test are multiplied by the same constant and the error variance is multiplied by the square of the constant.

Let  $D$  be a diagonal matrix with positive diagonal elements. Consider the model

$$\underline{\tilde{X}} = \underline{\tilde{\mu}} + \Lambda \underline{\tilde{f}} + \underline{\tilde{\epsilon}}$$

and premultiply it with  $D$  as

$$\begin{aligned} D\underline{\tilde{X}} &= D\underline{\tilde{\mu}} + D\Lambda \underline{\tilde{f}} + D\underline{\tilde{\epsilon}} \\ \text{or } \underline{\tilde{X}}^* &= \underline{\tilde{\mu}}^* + \Lambda^* \underline{\tilde{f}} + \underline{\tilde{\epsilon}}^* \end{aligned}$$

where  $\underline{\tilde{X}}^* = D\underline{\tilde{X}}$ ,  $\underline{\tilde{\mu}}^* = E(\underline{\tilde{X}}^*) = D\underline{\tilde{\mu}}$ ,  $\Lambda^* = D\Lambda$ ,  $\underline{\tilde{\epsilon}}^* = D\underline{\tilde{\epsilon}}$ ,  $Cov(\underline{\tilde{\epsilon}}^*) = D\Psi D$ .

Then

$$E(\underline{\tilde{X}}^* - \underline{\tilde{\mu}}^*)(\underline{\tilde{X}}^* - \underline{\tilde{\mu}}^*)' = \Lambda^* \Phi \Lambda^{*'} + \Psi^* = \Sigma^*$$

where  $\Sigma^* = D\Sigma D$ . If the identification conditions are  $\Phi = I$  and  $\Lambda'\Psi^{-1}\Lambda$  is diagonal, then  $\Lambda^*$  satisfies these conditions.



## Model with Orthogonal Factors

Now we restrict our attention to a model with orthogonal factors and develop the analysis further for this model.

For review, the model is now

$$\tilde{X} = \mu + \Lambda \tilde{f} + \epsilon$$

with  $E(\epsilon) = 0$ ,  $E(\tilde{f}) = 0$ ,  $E(\tilde{X}) = \mu$ ,  $Cov(\tilde{f}) = E(\tilde{f}\tilde{f}') = I$ ,  $Cov(\epsilon) = E(\epsilon\epsilon') = \Psi$  and the specific factors are uncorrelated with the common factors, i.e.,  $Cov(\epsilon_i, f_j) = 0$ ,  $i = 1, 2, \dots, p; j = 1, 2, \dots, m$ .

These assumptions are necessary for the unique estimation of the parameters.

The **communality** for variable  $i$  is  $\sum_{j=1}^m \lambda_{ij}^2$  and variance of trait  $i$  is

$$\sigma_i^2 = Var(X_i) = \sum_{j=1}^m \lambda_{ij}^2 + \psi_i.$$

The covariance between trait  $i$  and factor  $j$  is  $Cov(X_i, f_j) = \lambda_{ij}$ . So we can express

$$\Sigma = \Lambda\Lambda' + \Psi$$

where  $\Lambda$  is the matrix of factor loadings and  $\Psi = diag(\psi_1, \psi_2, \dots, \psi_p)$  is the diagonal matrix of specific variances.

It is worth noting that the model assumes the data are linear functions of unobservable common factors, so the linearity cannot be checked.

The covariance matrix  $\Sigma$ , having  $\frac{p(p+1)}{2}$  unique elements, is approximated by  $mp$  factor loadings in  $\Lambda$  and  $p$  specific variances. Ideally  $(mp + p) \ll \frac{p(p+1)}{2}$ . In case  $mp$  is too small, then  $(mp + p)$  parameters may not adequately describe  $\Sigma$ . In such conditions, the model may not be the right one, but we cannot reduce the data to a linear combination of factors.

It is important to note that the calculations in the original model and the model obtained by an orthogonal matrix remain the same. This is shown as follows.

With an  $m \times m$  orthogonal matrix  $D$ , the model

$$\tilde{X} = \mu + \Lambda \tilde{f} + \epsilon$$

is written as

$$\begin{aligned}\tilde{X} &= \mu + \Lambda DD' \tilde{f} + \epsilon \\ &= \mu + \Lambda^* \tilde{f}^* + \epsilon\end{aligned}$$

and does not change any calculations because  $DD' = I_m$  and multiplication by  $I$  to any matrix does not alter the original matrix.

Also,

$$\begin{aligned}E(\tilde{f}^*) &= E(D' \tilde{f}) = D' E(\tilde{f}) = 0 \\ Var(\tilde{f}^*) &= D' Var(\tilde{f}) D = D' I D = I \\ Cov(\tilde{f}^*, \epsilon) &= Cov(D' \tilde{f}, \epsilon) = D' Cov(\tilde{f}, \epsilon) = 0.\end{aligned}$$

Thus  $\tilde{f}^*$  satisfies all the assumptions and is an equally valid collection of common factors.

It is worth noting that there are different choices of orthogonal matrices that give rise to different models, but all these models satisfy the model assumptions.

Next we consider the estimation of parameters of a factor model. The parameters can be estimated by the principal component method, the maximum likelihood method, the principal factor method, etc.

## Principal Component Method of Estimation of Parameters

It may be recalled that the principal component method is used for dimension reduction and its methodology is based on the characteristic roots and characteristic vectors of covariance matrix. So, the principal component method should be understood in the context of the developments and details of principal component analysis.

Let the observations for the  $l^{th}$  subject on  $p$  variables is expressed in a  $p \times 1$  vector  $\tilde{X}_l = (X_{l1}, X_{l2}, \dots, X_{lp})'$  and its sample covariance matrix is

$$S = \frac{1}{n-1} \sum_{l=1}^n (\tilde{X}_l - \bar{\tilde{X}})(\tilde{X}_l - \bar{\tilde{X}})'$$

which is based on  $\underline{X}_1, \underline{X}_2, \dots, \underline{X}_n$ . We find the characteristic roots of  $S$  as  $\hat{\lambda}_1, \hat{\lambda}_2, \dots, \hat{\lambda}_p$  and the corresponding characteristic vectors as  $\hat{e}_1, \hat{e}_2, \dots, \hat{e}_p$ , respectively.

Using the spectral decomposition theorem over  $\Sigma$ , we can express

$$\Sigma = \sum_{i=1}^p \lambda_i \underline{e}_i \underline{e}_i' \approx \sum_{i=1}^m \lambda_i \underline{e}_i \underline{e}_i'$$

where the last  $(p - m)$  terms are ignored, similar to the case in principal component analysis, and  $\Sigma$  is approximated by first  $m$  terms. This can be further expressed in terms of  $\Lambda$  as

$$\begin{aligned} \Sigma &= \sum_{i=1}^m \lambda_i \underline{e}_i \underline{e}_i' \\ &= (\sqrt{\lambda_1} \underline{e}_1, \sqrt{\lambda_2} \underline{e}_2, \dots, \sqrt{\lambda_m} \underline{e}_m) \begin{pmatrix} \sqrt{\lambda_1} \underline{e}_1' \\ \sqrt{\lambda_2} \underline{e}_2' \\ \vdots \\ \sqrt{\lambda_m} \underline{e}_m' \end{pmatrix} \\ &= \Lambda \Lambda'. \end{aligned}$$

It is important to note that the maximum likelihood estimates of the characteristic roots and characteristic vectors in principal component analysis were obtained as the sample-based characteristic roots and characteristic vectors from the sample covariance matrix. In case the observations are standardized, the sample covariance matrix will be replaced by the sample correlation matrix.

Thus we obtain

$$\hat{\lambda}_{ij} = \hat{e}_{ji} \sqrt{\hat{\lambda}_j}$$

Since  $\Psi = \Sigma - \Lambda \Lambda'$ , so the diagonal elements of  $\Psi$ , are estimated by

$$\hat{\psi}_i = s_i^2 - \sum_{j=1}^m \hat{\lambda}_j \hat{e}_{ji}^2$$

where  $s_i^2$  is the sample variance for the  $i^{th}$  variable.

The factor loadings, which are the correlations between the factors and the variables, are obtained by

$$\hat{e}_i \sqrt{\hat{\lambda}_i}.$$

The interpretation of factor loadings is similar to that of the coefficients for principal components in principal component analysis. Some criterion for retaining the correlation coefficients is used, which may be arbitrary and can be developed within the framework of a test of hypothesis, as in the case of principal component analysis.

The communalities for the  $i^{th}$  variable are found from the factor loadings as

$$l_i^2 = \sum_{j=1}^m \hat{\lambda}_{ij}^2.$$

## Performance Evaluation

The communality for a given variable is interpreted as the proportion of variation in that variable explained by the required variables of interest. It is equivalent to the coefficient of determination  $R^2$  in the case of a multiple linear regression model in the sense that the model is fitted with these variables of interest.

The model performance can be assessed from the communalities. Values close to 1 are observed, indicating that the model explains most of the variation in these variables. In such a case, the model performs better for some variables than it does for others.

The sum of all communality values is the total communality value

$$\sum_{i=1}^p l_i^2 = \sum_{i=1}^m \lambda_i$$

The proportion of total variation explained by the selected factors is obtained as

$$\frac{\text{Total communality}}{\text{Number of variables}} = \frac{\sum_{i=1}^p l_i^2}{p}$$

which gives the percentage of variation explained by the model, and is considered as a measure of the overall assessment of the model performance. This proportion is the same as the proportion of variation explained by the  $m$  characteristic roots. The individual

communalities shed light on how the model performs for individual variables, and the total communality provides an overall assessment of performance.

As the data are standardized, so the variance of the standard data is equal to 1. The specific variances are computed by

$$\hat{\psi}_i = 1 - l_i^2$$

for all the variables.

Such a model provides an approximation to the correlation matrix, and the model's appropriateness can be assessed with the residuals obtained from the calculation

$$s_{ij} - \sum_{k=1}^m \lambda_{ik} \lambda_{jk}, \quad i \neq j = 1, 2, \dots, p$$

which is essentially the difference between  $R$  and  $\Lambda\Lambda'$ . Such residuals are expected to be close to 0 as far as possible for a good model fit.

The principal component method does not provide a test of lack of fit, and the decisions are based on the numbers the experimenter feels are suitable. Such a test is available for the maximum likelihood method.

When the number of factors increases, the estimated factor loadings do not change in the principal component method, but they change in the maximum likelihood method.

As the number of factors increases, the communalities increase and tend towards 1 whereas the specific variances decrease towards zero.

The diagonal elements of the sample covariance matrix  $S$  are equal to the diagonal elements of  $\hat{\Lambda}\hat{\Lambda}' + \hat{\Psi}$ , but the off-diagonal elements are not exactly reproduced due to variability in the data. Therefore, the number of factors is selected such that the off-diagonal elements of the residual matrix  $S - (\hat{\Lambda}\hat{\Lambda}' + \hat{\Psi})$  are small.

When we want to choose  $m$  to be as small as possible, then the residual for that model can be large. On the other hand, when  $m$  is chosen to be large, the residuals will become smaller, but the model will become more complex and less interpretable.

The sum of squared elements of the residual matrix is equal to the sum of the squared values of the eigenvalues left out of the matrix, i.e.,  $\sum_{j=m+1}^p \hat{\lambda}_j^2$ .

## Maximum Likelihood Estimation of Parameters

We assume the observations are sampled from a multivariate normal distribution with a continuous random vector. Data in the social sciences are often on a Likert scale, which is discrete and bounded, which poses a limitation on the maximum likelihood estimates that can be used.

Assuming the  $f_i$  and  $\epsilon_i$  to follow jointly normal, then  $x_i - \mu = \Lambda f_i + \epsilon_i$  are then normal and the likelihood function is

$$\begin{aligned} L(\mu, \Sigma) &= \frac{1}{(2\pi)^{\frac{Np}{2}} |\Sigma|^{\frac{N}{2}}} \exp \left[ -\frac{1}{2} \sum_{i=1}^N (X_i - \mu)' \Sigma^{-1} (X_i - \mu) \right] \\ &= \frac{1}{(2\pi)^{\frac{Np}{2}} |\Sigma|^{\frac{N}{2}}} \exp \left[ -\frac{1}{2} \text{tr} \left\{ \Sigma^{-1} \left( \sum_{i=1}^N (X_i - \bar{x})(X_i - \bar{x})' + n(\bar{x} - \mu)(\bar{x} - \mu)' \right) \right\} \right] \\ &= \frac{1}{(2\pi)^{\frac{(N-1)p}{2}} |\Sigma|^{\frac{N-1}{2}}} \exp \left[ -\frac{1}{2} \text{tr} \left\{ \Sigma^{-1} \left( \sum_{i=1}^N (X_i - \bar{x})(X_i - \bar{x})' \right) \right\} \right] \\ &\quad \times \frac{1}{(2\pi)^{\frac{p}{2}} |\Sigma|^{\frac{1}{2}}} \exp \left[ -\frac{n}{2} (\bar{x} - \mu)' \Sigma^{-1} (\bar{x} - \mu) \right]. \end{aligned}$$

The likelihood function depends on  $\Lambda$  and  $\Psi$  through  $\Sigma = \Lambda\Lambda' + \Psi$ . This model is still not well-defined because of the multiplicity of choices for  $\Lambda$  enabled by orthogonal transformations.

The log likelihood is maximized to obtain the estimates of  $\mu$ ,  $\Lambda$  and  $\Psi$ , but the solutions for these factor models are not unique. The equivalent models are obtained by rotation. For example, if  $\Lambda'\Psi^{-1}\Lambda$  is a diagonal matrix, then a unique solution can be obtained. In general, iterative computation methods are employed to obtain the estimates. So there are no closed-form maximum-likelihood estimates, and the process is computationally intensive and complex. Thus the maximum likelihood estimates of  $\Lambda$  and  $\Psi$  are by numerically maximizing the log likelihood and the maximum likelihood estimates of  $\mu$  is  $\bar{x}$  subject to  $\Lambda'\Psi^{-1}\Lambda$  being diagonal.

## Asymptotic Distribution of the Estimates

**Theorem:** If  $\Lambda$  and  $\Psi$  are identified by  $\Lambda'\Psi^{-1}\Lambda$  being diagonal, if the diagonal elements are different and ordered and if  $C \xrightarrow{Pr} \Psi + \Lambda\Lambda'$ , then  $\hat{\Psi} \xrightarrow{Pr} \Psi$  and  $\hat{\Lambda} \xrightarrow{Pr} \Lambda$  where  $C = \frac{A}{N}$  and  $A = \sum_{i=1}^N (\tilde{X}_i - \bar{\tilde{x}})(\tilde{X}_i - \bar{\tilde{x}})'$ .

A sufficient condition for  $C \xrightarrow{Pr} \Sigma$  is that  $(\tilde{f}', \tilde{\epsilon}')'$  has a distribution with finite second order moments.

The estimators  $\hat{\Lambda}$  and  $\hat{\Psi}$  are the solutions of

$$\hat{\Lambda}(\hat{\Lambda}'\hat{\Psi}^{-1}\hat{\Lambda} + I) = C\hat{\Psi}^{-1}\hat{\Lambda}$$

and

$$diag(\hat{\Lambda}\hat{\Lambda}' + \hat{\Psi}) = diag(C)$$

with  $\hat{\Lambda}'\hat{\Psi}^{-1}\hat{\Lambda}$  being diagonal.

## Goodness of Fit

The goodness of fit test is based on the comparison of the covariance matrix under the assumed model  $\Sigma = \Lambda\Lambda' + \Psi$  with the covariance matrix that can take any values.

We have considered a specific structure of  $\Lambda$  and  $\Psi$  (i.e., the diagonal elements of  $\Psi$  are equal to some specified values). A more general framework is to allow these elements to take any value.

The Bartlett corrected likelihood ratio test statistic

$$\chi^2 = \left[ n - 1 - \frac{2p + 4m - 5}{6} \right] \log \frac{|\hat{\Lambda}\hat{\Lambda}' + \hat{\Psi}|}{|\hat{\Sigma}|}.$$

The smaller values of  $\chi^2$  indicate that the factor model fits well and larger values of  $\chi^2$  indicate that the factor model does not fit well.

Under the null hypothesis

$H_0$  : Factor model adequately describes the relationships among the variables,

$$\chi^2 \sim \chi^2_{\frac{(p-m)^2 - p - m}{2}}.$$

So reject  $H_0$  if  $\chi^2 > \chi^2_{\frac{(p-m)^2 - p - m}{2}}$ .

## Factor Rotations

The motivation for factor rotations is due to the fact that the factor models are not unique. Recall that the models  $\tilde{X} = \tilde{\mu} + \Lambda \tilde{f} + \tilde{\epsilon}$  and  $\tilde{X} = \tilde{\mu} + \Lambda^* \tilde{f}^* + \tilde{\epsilon}$  where  $\Lambda^* = \Lambda D$  and  $\tilde{f}^* = D' \tilde{f}$  for some orthogonal matrix  $D$  such that  $DD' = I$  are equivalent. The choice of  $D$  corresponds to a particular factor rotation, and there are an infinite number of possible orthogonal matrices. So the aim here is to find an appropriate rotation, defined by an orthogonal matrix  $D$ , that yields easily interpretable factors.

This can be done by considering a scatter plot of factor loadings. The orthogonal matrix rotates the axes of this scatter plot and finds a rotation such that each of the  $p$  variables has a high loading on only one factor. Each rotation corresponds to a choice of orthogonal matrix  $D$ . Our aim is to rotate the axes to obtain a cleaner interpretation of the data and to define a new coordinate system so that, after rotation, the points fall close to the vertices of the new axes. For example, considering two factors, we want to find each of the plotted points at the four tips of the rotated axes. So the rotation accounts for the factor pattern plot and rotates the axes in such a way that the points fall close to the axes.

There are two types of rotation methods - orthogonal and oblique. In **orthogonal rotation** the rotated factors will remain uncorrelated whereas in **oblique rotation**, the resulting factors will be correlated. There are a number of different methods of rotation of each type. The most common orthogonal method is varimax rotation.

## Varimax Rotation

The varimax rotation is the most widely used rotation criterion. It seeks the rotated loadings that maximize the variance of the squared loadings for each factor with the objective to make some of the loadings as large as possible and the rest as small as possible in absolute value. The varimax rotation criterion encourages the detection of



factors each of which is related to few variables and discourages the detection of factors influencing all variables.

The varimax rotation scales the loadings by dividing them by the corresponding communality as

$$\tilde{\lambda}_{ij}^* = \frac{\hat{\lambda}_{ij}^*}{l_i}.$$

The varimax rotation finds the rotation that maximizes the  $\tilde{\lambda}_{ij}^*$ . The varimax procedure selects the rotation, or equivalently, the orthogonal matrix  $D$  that makes

$$V = \frac{1}{p} \sum_{j=1}^m \left[ \sum_{i=1}^p (\tilde{\lambda}_{ij}^*)^4 - \frac{1}{p} \left\{ \sum_{i=1}^p (\tilde{\lambda}_{ij}^*)^2 \right\}^2 \right]$$

as large as possible.

This is the sample variance of the standardized loadings for each factor summed over the  $m$  factors. After determining  $D$ , the loadings  $\tilde{\lambda}_{ij}^*$  are multiplied by  $l_i$  so that the original communalities are preserved. The simple interpretation from  $V$  is

$$V \propto \sum_{j=1}^m \left[ \begin{array}{c} \text{variance of the square of} \\ \text{(scaled) loadings for } j^{th} \text{ factor} \end{array} \right].$$

Maximizing  $V$  corresponds to the square of the loadings on each factor as much as possible. So we expect to find groups of large and negligible coefficients in any column of the rotated loadings matrix  $\hat{\Lambda}^*$ .

## Oblique Rotation

Orthogonal rotations are appropriate for factor models where the common factors are assumed to be independent. Oblique (non-orthogonal) rotations are also useful in the social sciences, where the factors considered are correlated. Oblique rotations are more complex than orthogonal rotations, as they can involve one or two coordinate systems, a system of primary axes, or a system of reference axes. The oblique rotation produces a pattern matrix that contains the factor loadings and a factor correlation matrix that includes the correlations between the factors. The commonly used oblique rotation techniques are Direct Oblimin and Promax.

The **Direct Oblimin technique** simplifies the structure and the mathematics of the output of the problem under consideration. The **Promax technique** involves raising the loadings to the fourth power, resulting in greater correlations among the factors and achieving a simple structure. The Promax technique is used in larger datasets because of its speed.

## Estimation of Factor Scores

A factor score can be considered to be a variable describing how much an individual would score on a factor.

Considering the model

$$\tilde{X}_i = \mu + \Lambda \tilde{f}_i + \tilde{\epsilon}_i, \quad i = 1, 2, \dots, n.$$

We aim to estimate the vectors of factor scores  $\tilde{f}_1, \tilde{f}_2, \dots, \tilde{f}_n$  for each observation using the ordinary least squares method, the weighted least squares method and the regression method.

### 1. Ordinary Least Squares Method

The vector of common factors for subject  $i$ , i.e.,  $\hat{\tilde{f}}_i$  is obtained by minimizing the sum of squares due to  $\tilde{\epsilon}$ 's as

$$\sum_{j=1}^p \epsilon_{ij}^2 = (\tilde{X}_i - \mu - \Lambda \tilde{f}_i)'(\tilde{X}_i - \mu - \Lambda \tilde{f}_i).$$

This is similar to obtaining  $\hat{\beta}$  from  $(y - X\beta)'(y - X\beta)$  in the usual multiple linear regression model  $y = X\beta + \epsilon$  and we get  $\hat{\tilde{\beta}} = (X'X)^{-1}X'y$ .

So we get the ordinary least squares estimates of  $\tilde{f}_i$  as

$$\hat{\tilde{f}}_i = (\Lambda'\Lambda)^{-1}\Lambda'(\tilde{X}_i - \mu).$$

In practice, we substitute  $\hat{\Lambda}$  and  $\bar{x}$  in place of  $\Lambda$  and  $\mu$ , respectively which are obtained from the sampled data and obtain

$$\hat{\tilde{f}}_i = (\hat{\Lambda}'\hat{\Lambda})^{-1}\hat{\Lambda}'(\tilde{X}_i - \bar{x}).$$

Next, using the principal component method with the unrotated factor loadings, this gives

$$\hat{f}_i = \left( \frac{1}{\sqrt{\hat{\lambda}_1}} \hat{e}'_1(\tilde{X}_i - \bar{\tilde{x}}), \frac{1}{\sqrt{\hat{\lambda}_2}} \hat{e}'_2(\tilde{X}_i - \bar{\tilde{x}}), \dots, \frac{1}{\sqrt{\hat{\lambda}_m}} \hat{e}'_m(\tilde{X}_i - \bar{\tilde{x}}) \right)'$$

where  $\hat{e}_1, \hat{e}_2, \dots, \hat{e}_m$  are the first  $m$  characteristic vectors corresponding to the characteristic roots  $\hat{\lambda}_1, \hat{\lambda}_2, \dots, \hat{\lambda}_m$  obtained from the sample data.

We observe that  $\hat{f}_i$ 's are approximately the same as those obtained in the principal component method with the unrotated factor loadings. Here

$$\begin{aligned} \frac{1}{n} \sum_{j=1}^n \hat{f}_j &= 0 \quad (\text{sample mean}) \\ \frac{1}{n-1} \sum_{j=1}^n \hat{f}_j \hat{f}_j' &= 1 \quad (\text{sample variance.}) \end{aligned}$$

## 2. Weighted Least Squares Estimation

The weighted least squares method gives more weight to variables in the estimation that have low specific variances. The factor model fits the data best for variables with low specific variances. Variables with low specific variances should provide more information about the true values of the specific factors.

Considering the model

$$\tilde{X}_i = \tilde{\mu} + \Lambda \tilde{f}_i + \tilde{\epsilon}_i,$$

the weighted least squares method provides  $\hat{\tilde{f}}_i$  that are obtained by minimizing

$$\sum_{j=1}^p \frac{\epsilon_{ij}^2}{\Psi_j} = (\tilde{X}_i - \tilde{\mu} - \Lambda \tilde{f}_i)' \Psi^{-1} (\tilde{X}_i - \tilde{\mu} - \Lambda \tilde{f}_i).$$

This is again similar to obtaining the generalized or weighted least squares estimator  $\hat{\tilde{\beta}}$  in the usual multiple linear regression model  $\tilde{y} = \tilde{X}\tilde{\beta} + \tilde{\epsilon}$  with  $Cov(\tilde{\epsilon}) = \Psi$  and we get  $\hat{\tilde{\beta}} = (\tilde{X}'\Psi^{-1}\tilde{X})^{-1}\tilde{X}'\Psi^{-1}\tilde{y}$ .

So we get the generalized least squares estimator of  $\tilde{f}_i$  as

$$\hat{\tilde{f}}_i = (\Lambda'\Psi^{-1}\Lambda)^{-1}\Lambda'\Psi^{-1}(\tilde{X}_i - \tilde{\mu}).$$

In practice, we substitute  $\hat{\Lambda}$  and  $\bar{x}$  in place of  $\Lambda$  and  $\mu$ , respectively and obtain

$$\hat{f}_i = (\hat{\Lambda}'\hat{\Psi}^{-1}\hat{\Lambda})^{-1}\hat{\Lambda}'\hat{\Psi}^{-1}(X_i - \bar{x}).$$

### 3. Regression Method of Estimation

This method is used for maximum likelihood estimates of factor loadings considering the joint distribution of  $X_i$  and  $f_i$  as

$$\begin{pmatrix} X_i \\ f_i \end{pmatrix} \sim N \left[ \begin{pmatrix} \mu \\ 0 \end{pmatrix}, \begin{pmatrix} \Lambda\Lambda' + \Psi & \Lambda \\ \Lambda' & I \end{pmatrix} \right].$$

Next obtain the conditional expectation of the common factor score  $f_i$  given  $X_i$  as

$$E(f_i|X_i) = \Lambda'(\Lambda\Lambda' + \Psi)^{-1}(X_i - \mu).$$

In practice, we replace  $\Lambda$  and  $\mu$  by  $\hat{\Lambda}$  and  $\bar{x}$  and obtain

$$\hat{f}_i = \hat{\Lambda}'(\hat{\Lambda}\hat{\Lambda}' + \hat{\Psi})^{-1}(X_i - \bar{x}).$$

### Estimation for Fixed Models

Let  $x_\alpha = (x_{1\alpha}, x_{2\alpha}, \dots, x_{p\alpha})'$  be an observation on  $X_\alpha$ ,  $\alpha = 1, 2, \dots, N$ , given by

$$X_\alpha = \mu + \Lambda f_\alpha + \xi_\alpha$$

with  $f_\alpha$  being nonstochastic vector satisfying  $\sum_{\alpha=1}^N f_\alpha = 0$ .

The likelihood function is

$$L = \frac{1}{[(2\pi)^p \prod_{i=1}^p \psi_{ii}]^{\frac{N}{2}}} \prod_{i=1}^p \exp \left[ -\frac{1}{2} \sum_{\alpha=1}^N \frac{(x_{i\alpha} - \mu_i - \sum_{j=1}^m \lambda_{ij} f_{j\alpha})^2}{\psi_{ii}} \right].$$

The likelihood function does not have a maximum. To show this fact, let  $\mu_i = 0$ ,  $\lambda_{11} = 1$ ,  $\lambda_{1j} = 0$  ( $j \neq 1$ ),  $f_{1\alpha} = x_{1\alpha}$ . Then  $x_{1\alpha} - \mu_1 - \sum_{j=1}^m \lambda_{1j} f_{j\alpha} = 0$  and  $\psi_{11}$  does not appear in the exponent but appears only in the constant. As  $\psi_{11} \rightarrow 0$ ,  $L \rightarrow \infty$ . Thus, the likelihood function does not have a maximum, and the maximum-likelihood estimators do not exist.

Under such a condition, one possibility is to use the maximum likelihood method appropriate for random factors. The sample covariance matrix under normality has a noncentral Wishart distribution, depending on  $\Psi$ ,  $\Lambda\Phi\Lambda'$ , and  $N - 1$ .