

- Then, the distribution of

$$N_p(\mu, \Sigma = \mathbf{L}\mathbf{L}' + \Psi) \quad \mathbf{X} = \mu + \mathbf{L}\mathbf{F} + \varepsilon \sim N_p(0, \Psi)$$

is multivariate normal and its likelihood is

$$\begin{aligned} L(\mu, \Sigma) &= (2\pi)^{-\frac{np}{2}} |\Sigma|^{-\frac{n}{2}} e^{-\frac{1}{2} \text{tr} \left[\Sigma^{-1} \left(\sum_{j=1}^n (\mathbf{x}_j - \bar{\mathbf{x}})(\mathbf{x}_j - \bar{\mathbf{x}})' + n(\bar{\mathbf{x}} - \mu)(\bar{\mathbf{x}} - \mu)' \right) \right]} \\ &= (2\pi)^{-\frac{(n-1)p}{2}} |\Sigma|^{-\frac{(n-1)}{2}} e^{-\frac{1}{2} \text{tr} \left[\Sigma^{-1} \left(\sum_{j=1}^n (\mathbf{x}_j - \bar{\mathbf{x}})(\mathbf{x}_j - \bar{\mathbf{x}})' \right) \right]} \\ &\quad \times (2\pi)^{-\frac{p}{2}} |\Sigma|^{-\frac{1}{2}} e^{-\frac{n}{2} (\bar{\mathbf{x}} - \mu)' \Sigma^{-1} (\bar{\mathbf{x}} - \mu)} \end{aligned}$$

cf. the
Σ in
principal component approach

where $\Sigma = \mathbf{L}\mathbf{L}' + \Psi \Rightarrow$ parameters, $m, \mathbf{L}, \Psi, \mu$.

- To make $\mathbf{L}$ well defined, we can impose the computationally convenient uniqueness condition

$$(\Psi^{-1/2} \mathbf{L})' (\Psi^{-1/2} \mathbf{L}) = \mathbf{L}' \Psi^{-1} \mathbf{L} = \Delta \quad \text{a diagonal matrix} \Rightarrow \begin{bmatrix} \times & 0 \\ 0 & \times \end{bmatrix} \quad \text{how many constraints?} \quad \# \text{ of } 0\text{'s} = \frac{m^2 - m}{2}$$

(Q: dimension reduction of the parameter space caused by the condition $= \frac{m(m-1)}{2}$)

- Maximizing the likelihood is equivalent to minimizing

$$\ln |\Sigma| - \ln |\mathbf{S}_n| + \text{tr}(\Sigma^{-1} \mathbf{S}_n) - p \quad (\diamond) \quad \leftarrow \text{MLE of } \Sigma \quad \text{(for details, see Supplement 9A).}$$

subject to $\mathbf{L}' \Psi^{-1} \mathbf{L} = \Delta$, a diagonal matrix.

- Note: the function take the value zero if $\Sigma = \mathbf{L}\mathbf{L}' + \Psi$ equals $\mathbf{S}_n$
- The MLE of $\mathbf{L}$ and Ψ can be obtained by numerical maximization

- Result 9.1.** Let $\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_n$ be a random sample from $N_p(\mu, \Sigma)$, where $\Sigma = \mathbf{L}\mathbf{L}' + \Psi$ is the covariance matrix for the m common factor model. The maximum likelihood estimators $\hat{\mathbf{L}}, \hat{\Psi}$, and $\hat{\mu} = \bar{\mathbf{x}}$ maximize (or minimize) $(\diamond)$ subject to $\hat{\mathbf{L}}' \hat{\Psi}^{-1} \hat{\mathbf{L}}$ being diagonal.

The maximum likelihood estimates of the communalities are

$$\hat{h}_i^2 = \hat{\ell}_{i1}^2 + \hat{\ell}_{i2}^2 + \dots + \hat{\ell}_{im}^2 \quad \text{for } i = 1, 2, \dots, p$$

so

$$\left(\begin{array}{c} \text{Proportion of total sample} \\ \text{variance due to } j\text{th factor} \end{array} \right) = \frac{\hat{\ell}_{1j}^2 + \hat{\ell}_{2j}^2 + \dots + \hat{\ell}_{pj}^2}{s_{11} + s_{22} + \dots + s_{pp}} \quad \begin{array}{l} \theta \uparrow \\ \hat{\theta}_{MLE} \end{array} \quad \begin{array}{l} \eta = g(\theta) \\ \hat{\eta}_{MLE} = g(\hat{\theta}_{MLE}) \end{array}$$

- If the variables are standardized so that $\mathbf{Z} = \mathbf{V}^{-1/2}(\mathbf{X} - \mu)$, then the covariance matrix ρ of $\mathbf{Z}$ has the representation

correlation matrix of $\mathbf{X}$

$$\rho = \mathbf{V}^{-1/2} \Sigma \mathbf{V}^{-1/2} = (\mathbf{V}^{-1/2} \mathbf{L}) (\mathbf{V}^{-1/2} \mathbf{L})' + \mathbf{V}^{-1/2} \Psi \mathbf{V}^{-1/2}$$

Thus, ρ has a factorization with loading matrix $\mathbf{L}_z = \mathbf{V}^{-1/2} \mathbf{L}$ and

specific variance matrix $\Psi_z = \mathbf{V}^{-1/2} \Psi \mathbf{V}^{-1/2}$

- By the invariance property of MLE, the MLE of ρ is

$$\begin{aligned} \hat{\rho} &= (\hat{\mathbf{V}}^{-1/2} \hat{\mathbf{L}}) (\hat{\mathbf{V}}^{-1/2} \hat{\mathbf{L}})' + \hat{\mathbf{V}}^{-1/2} \hat{\Psi} \hat{\mathbf{V}}^{-1/2} \\ &= \hat{\mathbf{L}}_z \hat{\mathbf{L}}_z' + \hat{\Psi}_z \end{aligned}$$

where $\hat{\mathbf{V}}^{-1/2}$ and $\hat{\mathbf{L}}$ are the maximum likelihood estimators of $\mathbf{V}^{-1/2}$ and $\mathbf{L}$.

p. 5-14

◆ Note: There is usually no relationship between the principal components of the sample correlation matrix and the sample covariance matrix. For maximum likelihood factor analysis, however, the results of analyzing either matrix are essentially equivalent (this is not true of principal factor analysis)

↳ Note: not $\hat{L}_Z = \hat{L}$

■ The MLE method could produce very different results when $m \rightarrow m+1$

■ The MLE approach can also experience difficulties with Heywood cases

■ comparison of the PC and MLE approaches (Example: stock-price data)

Variable	Maximum likelihood <i>not orthogonal</i>			Principal components <i>orthogonal</i>		
	Estimated factor loadings F_1	F_2	Specific variances $\hat{\psi}_i = 1 - \hat{h}_i^2$	Estimated factor loadings F_1	F_2	Specific variances $\tilde{\psi}_i = 1 - \tilde{h}_i^2$
1. J P Morgan	.115	.755	.42	.732	-.437	.27
2. Citibank	.322	.788	.27	.831	-.280	.23
3. Wells Fargo	.182	.652	.54	.726	-.374	.33
4. Royal Dutch Shell	1.000	-.000	.00	.605	.694	.15
5. Texaco	.683	-.032	.53	.563	.719	.17
Cumulative proportion of total (standardized) sample variance explained	.323		$\frac{\sum 10}{5} = 1 - 0.647$ $0.647 - 0.323 = 0.324$	.487		$0.769 - 0.487 = 0.282$.769

- orthogonality
- estimated value of loading (very different $\leftarrow$ $\therefore$ different constraints)
- total variance explained.

◆ residual matrix $(\mathbf{I} - \mathbf{L}\mathbf{L}' - \mathbf{\Psi})$

p. 5-15

MLE

0	.001	-.002	.000	.052
.001	0	.002	.000	-.033
-.002	.002	0	.000	.001
.000	.000	.000	0	.000
.052	-.033	.001	.000	0

PC

0	-.099	-.185	-.025	.056
-.099	0	-.134	.014	-.054
-.185	-.134	0	.003	.006
-.025	.014	.003	0	-.156
.056	-.054	.006	-.156	0

- A large sample test for the number of common factors (i.e., m)

➤ Recall: choosing m might be done by examining the total variance explained

➤ Assume the common factors and the specific factors are jointly normally distributed

➤ null and alternative hypotheses

$$H_0: \mathbf{\Sigma} = \mathbf{L} \mathbf{L}' + \mathbf{\Psi} \quad \text{subject to } \mathbf{L}'\mathbf{\Psi}^{-1}\mathbf{L} = \mathbf{\Delta}$$

$H_1: \mathbf{\Sigma}$ any other positive definite matrix

- under alternative, $\hat{\mu} = \bar{\mathbf{x}}$ and $\hat{\Sigma} = ((n-1)/n)\mathbf{S} = \mathbf{S}_n$, the maximized likelihood is proportional to $|\mathbf{S}_n|^{-n/2} e^{-np/2}$

- under null, $\hat{\mu} = \bar{\mathbf{x}}$ and $\hat{\Sigma} = \hat{\mathbf{L}}\hat{\mathbf{L}}' + \hat{\mathbf{\Psi}}$, where $\hat{\mathbf{L}}$ and $\hat{\mathbf{\Psi}}$ are the MLE of $\mathbf{L}$ and $\mathbf{\Psi}$, the maximized likelihood is proportional to

$$|\hat{\mathbf{L}}\hat{\mathbf{L}}' + \hat{\mathbf{\Psi}}|^{-n/2} \exp\left(-\frac{1}{2}n \text{tr}[(\hat{\mathbf{L}}\hat{\mathbf{L}}' + \hat{\mathbf{\Psi}})^{-1}\mathbf{S}_n]\right) = p \text{ (supplement 9A)}$$

➤ likelihood ratio statistic

$$-2 \ln \Lambda = -2 \ln \left[\frac{\text{maximized likelihood under } H_0}{\text{maximized likelihood}} \right] = n \ln \left(\frac{|\hat{\Sigma}|}{|\mathbf{S}_n|} \right)$$

with degrees of freedom, $\sim \chi^2_{\dim(H_1) - \dim(H_0)}$

of parameters in a covariance matrix $\hat{\Sigma} \hat{\Sigma}' + \Psi$

of parameters in L, Ψ

$$v - v_0 = \left[\frac{1}{2}p(p+1) \right] - \left[p(m+1) - \frac{1}{2}m(m-1) \right]$$

$$= \frac{1}{2}[(p-m)^2 - p - m]$$

→ # of constraints in $L'\Psi^{-1}L = \Delta$

- Bartlett (1954) has shown that the chi-square approximation can be improved by replacing n with $(n - 1 - (2p + 4m + 5)/6)$
- Using Bartlett's correction, we reject H_0 at the α significant level if

$$(n - 1 - (2p + 4m + 5)/6) \ln \frac{|\hat{L}\hat{L}' + \hat{\Psi}|}{|\mathbf{S}_n|} > \chi^2_{[(p-m)^2 - p - m]/2}(\alpha)$$

> 0

- Note: because the number of degrees of freedom must be positive, it follows that

$$m < \frac{1}{2}(2p + 1 - \sqrt{8p + 1})$$

in order to apply the test

❖ Reading: Textbook, 9.3, Supplement 9A

• factor rotation

➤ Recall

- a solution of factor loading is not unique
- all factor loadings obtained from the initial loadings by an orthogonal transformation have the same ability to reproduce the covariance/correlation matrix

add constraints

If $\hat{L}$ is the $p \times m$ matrix of estimated factor loadings obtained by any method (principal component, maximum likelihood, and so forth) then

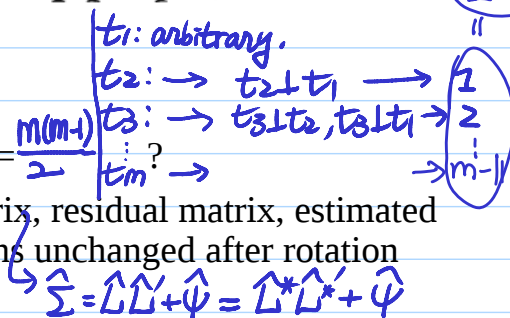
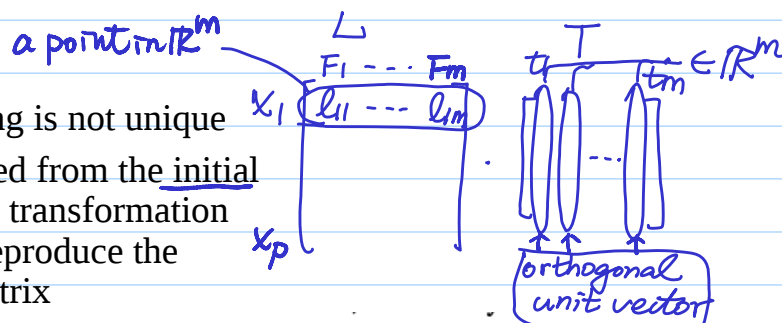
$$\hat{L}^* = \hat{L}T, \quad \text{where } TT' = T'T = I \quad \hat{L}\hat{L}' = \hat{L}^*\hat{L}^{*'} \quad \frac{m(m-1)}{2}$$

is a $p \times m$ matrix of "rotated" loadings.

- **Q:** why is it called "rotation"?
- **Q:** dimension of all rotated factor loadings = $\frac{m(m-1)}{2}$?
- Note: estimated covariance/correlation matrix, residual matrix, estimated specific variances and communalities remains unchanged after rotation

➤ **Q:** Why need rotation?

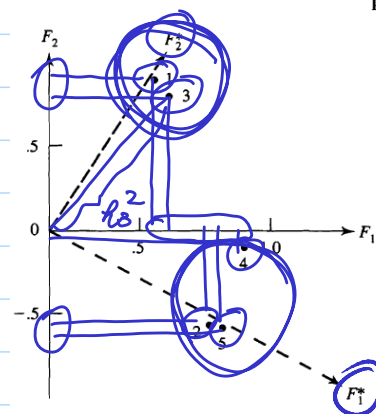
- since the original loading may not be readily interpretable (e.g., some factors may have several large loadings and some loadings may be positive while other may be negative) it's usual practice to rotate them until a "simple structure" is achieved *often seen in PC approach.*



p. 5-18

➤ **Q:** What is a “simple structure”?

- Ideally, we should like to see a pattern of loadings such that each variable loads highly on a *single* factor and has small to moderate loading on the remaining factors
- When there are two factors, we can plot the loadings and visually rotate the data to find an interpretable solution
- For more than two factors, some automation is necessary, which should make:
 - ◆ each factor have a few high positive loadings and the remainder be small
 - ◆ pairs of factors have few large loading in common. A partition into mutually exclusive groups would be desirable



Variable	Estimated factor loadings		Rotated estimated factor loadings		Communalities $\hat{h}_i^2$
	F_1	F_2	F_1^*	F_2^*	
1. Taste	.56	.82	.02	.99	.98
2. Good buy for money	.78	-.52	.94	-.01	.88
3. Flavor	.65	.75	.13	.98	.98
4. Suitable for snack	.94	-.10	.84	.43	.89
5. Provides lots of energy	.80	-.54	.97	-.02	.93
Cumulative proportion of total (standardized) sample variance explained	.571	.932	.507	.932	

Handwritten notes: "same" next to the communalities column, and "different" with an arrow pointing from the cumulative variance explained for F1 and F2.

p. 5-19

➤ varimax criterion

- Define $\tilde{\ell}_{ij}^* = \hat{\ell}_{ij}^* / \hat{h}_i$ (the scaling has the effect of giving variables with small communalities relatively more weight in rotation)
- varimax procedure selects the orthogonal transformation $\mathbf{T}$ that makes

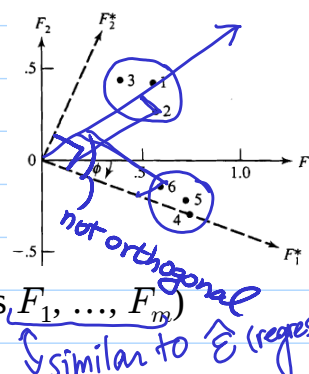
$$V = \frac{1}{p} \sum_{j=1}^m \left[\left(\frac{1}{p} \sum_{i=1}^p \tilde{\ell}_{ij}^{*4} \right) - \left(\frac{1}{p} \sum_{i=1}^p \tilde{\ell}_{ij}^{*2} \right)^2 \right]$$

as large as possible. $= \sum_j \left\{ \frac{1}{p} (\sum \ell^2)^2 - \left[\frac{1}{p} (\sum \ell^2) \right]^2 \right\}$

- interpretation of the varimax criterion: $\frac{EY^2}{Y=X^2} = (EY)^2$
 $V \propto \sum_{j=1}^m \left(\text{variance of squares of (scaled) loadings for } j\text{th factor} \right)$

- ◆ maximizing V corresponds to “spreading out” the squares of the loadings on each factor as much as possible

➤ Remark: In some cases, even orthogonal rotations do not provide an easy interpretation of the solution. It is possible to allow non-orthogonal rotations, called *oblique* rotation. This allows for possible simplicity at the expense of losing the orthogonality of the factors $\text{cov}(F^*) \neq I$.



- factor scores (predict values of the unobserved random factors $F_1, \dots, F_m$)
 - (c.f.) the scores in principal component analysis
- Handwritten note: "similar to $\hat{e}$ (regression)" with an arrow pointing to the factor scores.

➤ weighted least squares method

- Suppose first that the mean vector μ , the factor loadings L , and the specific variance Ψ are known for the factor model

$$\underset{(p \times 1)}{X} - \underset{(p \times 1)}{\mu} = \underset{(p \times m)}{L} \underset{(m \times 1)}{F} + \underset{(p \times 1)}{\epsilon} \quad \text{minimize } (Y - X\beta)' \Psi^{-1} (Y - X\beta)$$

- The sum of the squares of the errors, weighted by the reciprocal of their variances, is

$$\sum_{i=1}^p \frac{\epsilon_i^2}{\psi_i} = \epsilon' \Psi^{-1} \epsilon = (x - \mu - Lf)' \Psi^{-1} (x - \mu - Lf) \quad (*)$$

$$= [\psi^{-1/2}(x - \mu - Lf)]' [\psi^{-1/2}(x - \mu - Lf)]$$

- choosing the estimates $\hat{f}$ of f to minimize (*). The solution is

$$\hat{f} = (L' \Psi^{-1} L)^{-1} L' \Psi^{-1} (x - \mu) = \left[(\psi^{-1/2} L)' (\psi^{-1/2} L) \right]^{-1} (\psi^{-1/2} L)' \psi^{-1/2} (x - \mu)$$

- take the estimates $\hat{L}$, $\hat{\Psi}$, and $\hat{\mu} = \bar{x}$ as the true values and obtain the factor scores for the j th case as

$$\hat{f}_j = (\hat{L}' \hat{\Psi}^{-1} \hat{L})^{-1} \hat{L}' \hat{\Psi}^{-1} (x_j - \bar{x})$$

- When MLE method is used, $\hat{L}' \hat{\Psi}^{-1} \hat{L} = \hat{\Delta}$ is a diagonal matrix

$$\hat{f}_j = \hat{\Delta}^{-1} \hat{L}' \hat{\Psi}^{-1} (x_j - \bar{x})$$

- if the correlation matrix is factored

$$\hat{f}_j = (\hat{L}_z' \hat{\Psi}_z^{-1} \hat{L}_z)^{-1} \hat{L}_z' \hat{\Psi}_z^{-1} z_j = \hat{\Delta}_z^{-1} \hat{L}_z' \hat{\Psi}_z^{-1} z_j$$

where $z_j = D^{-1/2}(x_j - \bar{x})$