

Chapter 1

Principal Components Analysis

June 24, 2013

G Tutz - preliminary

1.1	Principal Components for Random Variables	3
1.1.1	Basic Concept	3
1.1.2	Obtaining Solutions	4
1.1.3	Variation and Explained Variation	6
1.1.4	Some Geometry and the Normal Distribution	7
1.2	Principal Components for Observations	7
1.3	Estimation	9

The objective of principal components analysis is to transform a set of p variables $x_1, \dots, x_p$ into r new variables $z_1, \dots, z_r$ which account for most of the variation of the original variables. The number r of the so-called principal components $z_1, \dots, z_r$ is assumed to be small relative to p . Classical principal components analysis is based on linear transformations of the original variables.

1.1 Principal Components for Random Variables

Principal components analysis may be formulated for random variables or observations in a rather similar way. We will first consider principal components for random variables

1.1.1 Basic Concept

For the random variables form one assumes that $\mathbf{x}^T = (x_1, \dots, x_p)$ is a vector of random variables with mean $\boldsymbol{\mu}^T = (\mu_1, \dots, \mu_p)$ and covariance $\text{cov}(\mathbf{x}) = \boldsymbol{\Sigma}$.

The first principal component is determined by a weight vector $\boldsymbol{\alpha}_1$ chosen such that

$$z_1 = \boldsymbol{\alpha}_1^T \mathbf{x} = \alpha_{11}x_1 + \dots + \alpha_{1p}x_p$$

has maximal variance, that is,

$$\text{var}(z_1) = \boldsymbol{\alpha}_1^T \boldsymbol{\Sigma} \boldsymbol{\alpha}_1 \rightarrow \max_{\boldsymbol{\alpha}_1}.$$

under the constraint $\|\boldsymbol{\alpha}_1\| = 1$. The constraint is necessary since the variance could be increased without limit by increasing the components of $\boldsymbol{\alpha}_1$.

For the second principal component $z_2 = \boldsymbol{\alpha}_2^T \mathbf{x}$ one postulates

$$\text{var}(z_2) \rightarrow \max_{\boldsymbol{\alpha}_2}$$

with the constraints $\|\boldsymbol{\alpha}_2\| = 1$ and $\boldsymbol{\alpha}_1^T \boldsymbol{\alpha}_2 = 0$. The latter constraint is equivalent to postulating $\text{cov}(z_1, z_2) = \text{cov}(\boldsymbol{\alpha}_1^T \mathbf{x}, \boldsymbol{\alpha}_2^T \mathbf{x}) = \boldsymbol{\alpha}_1^T \boldsymbol{\Sigma} \boldsymbol{\alpha}_2 = 0$. The further principal components are obtained by looking for weights which maximize the variance under the restriction that the weight is orthogonal to the weights of the previous principal components.

In summary the objective is to find weights $\boldsymbol{\alpha}_1, \dots, \boldsymbol{\alpha}_r$ and corresponding linear combinations

$$\begin{aligned} z_1 &= \boldsymbol{\alpha}_1^T \mathbf{x} = \alpha_{11}x_1 + \dots + \alpha_{1p}x_p \\ &\vdots \\ z_r &= \boldsymbol{\alpha}_r^T \mathbf{x} = \alpha_{r1}x_1 + \dots + \alpha_{rp}x_p \end{aligned}$$

such that

$$\text{var}(z_j) = \boldsymbol{\alpha}_j^T \boldsymbol{\Sigma} \boldsymbol{\alpha}_j \rightarrow \max_{\boldsymbol{\alpha}_j}, \quad j = 1, \dots, p, \quad (1.1)$$

with the side constraints $\|\boldsymbol{\alpha}_j\| = 1, \boldsymbol{\alpha}_j^T \boldsymbol{\alpha}_s = 0, s = 1, \dots, j-1$. The second side constraint is equivalent to postulating that $\text{cov}(z_j, z_s) = 0, s = 1, \dots, j-1$.

Principal Components by Maximization of Variance

Find weights $\boldsymbol{\alpha}_1, \dots, \boldsymbol{\alpha}_p$ such that for $z_j = \boldsymbol{\alpha}_j^T \mathbf{x}$

$$\text{var}(z_j) = \boldsymbol{\alpha}_j^T \boldsymbol{\Sigma} \boldsymbol{\alpha}_j \rightarrow \max_{\boldsymbol{\alpha}_j}$$

with side constraints $\|\boldsymbol{\alpha}_j\| = 1, \boldsymbol{\alpha}_j^T \boldsymbol{\alpha}_s = 0, s = 1, \dots, j-1$.

1.1.2 Obtaining Solutions

For the first principal component, the problem is to maximize $\text{var}(\boldsymbol{\alpha}_1^T \mathbf{x}) = \boldsymbol{\alpha}_1^T \boldsymbol{\Sigma} \boldsymbol{\alpha}_1$ subject to $\boldsymbol{\alpha}_1^T \boldsymbol{\alpha}_1 = 1$. Using Lagrange multiplier λ one considers

$$\varphi_1(\boldsymbol{\alpha}_1) = \boldsymbol{\alpha}_1^T \boldsymbol{\Sigma} \boldsymbol{\alpha}_1 - \lambda(\boldsymbol{\alpha}_1^T \boldsymbol{\alpha}_1)$$

The derivatives of $\varphi(\boldsymbol{\alpha}_1)$

$$\begin{aligned} \frac{\partial \varphi_1}{\partial \boldsymbol{\alpha}} &= 2\boldsymbol{\Sigma} \boldsymbol{\alpha}_1 - 2\lambda \boldsymbol{\alpha}_1 \\ \frac{\partial \varphi_1}{\partial \lambda} &= \boldsymbol{\alpha}_1^T \boldsymbol{\alpha}_1 \end{aligned}$$

yield the equations

$$\Sigma \alpha_1 = \lambda \alpha_1, \quad \alpha_1^T \alpha_1 = 1$$

which represent an eigenvalue problem. Thus the eigenvector α_1 which corresponds to the largest eigenvalue λ_1 is a solution of the maximization problem.

Consider now maximization of $\text{var}(\alpha_2^T x) = \alpha_2^T \Sigma \alpha_2$ subject to constraints $\|\alpha_1\| = 1$, $\alpha_2^T \alpha_1 = 0$. one considers the function with Lagrange multipliers for the constraints

$$\varphi_2(\alpha_2) = \alpha_2^T \Sigma \alpha_2 + \lambda(\alpha_1^T \alpha_2 - 1) + \gamma(\alpha_1^T \alpha_2)$$

The derivatives $\partial \varphi_2 / \partial \alpha_2$, $\partial \varphi_2 / \partial \lambda$, $\partial \varphi_2 / \partial \gamma$ yield the equations

$$\begin{aligned} 2\Sigma \alpha_2 + 2\lambda \alpha_2 + \gamma \alpha_1 &= 0, \\ \alpha_1^T \alpha_2 &= 1, \\ \alpha_1^T \alpha_2 &= 0. \end{aligned}$$

Multiplication of the first equation with α_1^T (from the left side) yields $2\alpha_1^T \Sigma \alpha_2 + \gamma = 0$. Since α_1 is an eigenvector of Σ one has $\alpha_1^T \Sigma \alpha_2 = \alpha_2^T \Sigma \alpha_1 = \lambda_1 \alpha_2^T \alpha_1 = 0$ yielding $\gamma = 0$. Therefore the first equation has the form

$$\Sigma \alpha_2 = \lambda \alpha_2 \quad \text{subject to} \quad \alpha_2^T \alpha_1 = 0$$

The solution is the eigenvector α_2 for the second largest eigenvalue λ_2 .

Straightforward derivation shows that starting from eigenvector solutions $\alpha_1, \dots, \alpha_s$ maximization of $\text{var}(\alpha^T x)$ subject to $\alpha^T \alpha = 1$, $\alpha^T \alpha_j = 0$, $j = 1, \dots, s$ yields the eigenvector α_{s+1} corresponding to the next largest eigenvalue λ_{s+1} .

In summary the solutions of the maximization problem are the eigenvectors $\alpha_1, \dots, \alpha_p$ that correspond to eigenvalues $\lambda_1 \geq \dots \geq \lambda_p$. The spectral decomposition theorem yields that the symmetric covariance matrix Σ may be written as

$$\Sigma = P \Lambda P^T$$

where the columns of the orthogonal matrix $P = (\alpha_1, \dots, \alpha_p)$ are the eigenvectors $\alpha_1, \dots, \alpha_p$ of Σ and $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_p)$ is a diagonal matrix which has the eigenvalues $\lambda_1 \geq \dots \geq \lambda_p$ in the diagonal.

For positive definite covariance matrix Σ all the eigenvalues $\lambda_1, \dots, \lambda_p$ are positive. One obtains the principal components $z_i = \alpha_i^T x$, $i = 1, \dots, p$ in vector form by

$$z = P^T x$$

where $z^T = (z_1, \dots, z_p)$. The principal components represent uncorrelated linear combinations of the variables. One obtains

$$\text{cov}(z) = P^T \Sigma P = \Lambda$$

and therefore $\text{var}(z_i) = \lambda_i$, $\text{cov}(z_i, z_j) = 0$, $i \neq j$.

Principal Components

The weights of principal components $\alpha_1, \dots, \alpha_p$ are found as the columns of the spectral decomposition

$$\Sigma = P\Lambda P^T,$$

where $P = (\alpha_1, \dots, \alpha_p)$, $\Sigma = \text{diag}(\lambda_1, \dots, \lambda_p)$.

For the vector of principal components $z = P^T x$ one has

$$\text{cov}(z) = \Lambda$$

and for the variation of x and z one has

$$\text{tr}(\text{cov}(x)) = \text{tr}(\text{cov}(z))$$

$$|\text{cov}(x)| = |\text{cov}(z)|$$

1.1.3 Variation and Explained Variation

Variation of random vectors may be measured in several ways. A simple measure of the variation in vector x is the *total variation*. For correlated variables $x_1, \dots, x_p$ one obtains

$$\text{tvar}(x) = \sum_{i=1}^p \text{var}(x_i) = \text{tr}(\Sigma)$$

By considering

$$\text{tr}(\Sigma) = \text{tr}(P\Lambda P^T) = \text{tr}(\Lambda P^T P) = \text{tr}(\Lambda) = \sum_{i=1}^p \text{var}(y_i) = \text{tvar}(y)$$

one obtains that the total variation of x is the same as the total variation of the principal components y , that is,

$$\sum_{i=1}^p \text{var}(x_i) = \sum_{i=1}^p \text{var}(y_i).$$

A more general measure of variation in vector y is the *generalized variance* given as determined $|\text{cov}(x)|$. Comparison of the generalized variance of x and the principal components y yields

$$|\text{cov}(x)| = |\Sigma| = |P\Lambda P^T| = |P||\Lambda||P^T| = |\Lambda| = |\text{cov}(y)|$$

since P is an orthogonal matrix and therefore $|P||P^T| = 1$. In summary one has

$$|\text{cov}(x)| = |\text{cov}(y)|.$$

One may wonder if all principal components are necessary. Their ordering $\text{var}(y_1) = \lambda_1 \geq \dots \geq \text{var}(y_p) = \lambda_p$ suggests to consider which part of the variation is explained by the first r principal components $z_1, \dots, z_r$.

A simple measure for the explained variation is the *proportion of total variation*

$$t(r) = \frac{\sum_{i=1}^r \text{var}(y_i)}{\sum_{i=1}^p \text{var}(y_i)} = \frac{\lambda_1 + \dots + \lambda_r}{\lambda_1 + \dots + \lambda_p}.$$

Thus, if, for example, $t(2) = 0.8$ 80 percent of the total variation is explained by the first principal component.

1.1.4 Some Geometry and the Normal Distribution

The components of a point $\mathbf{x}^T = (x_1, \dots, x_p)$ from $\mathbb{R}^p$ can be seen as the coordinates when the $\mathbb{R}^p$ is spanned by the unit vectors $\mathbf{e}_1^T = (1, 0, \dots, 0), \dots, \mathbf{e}_p^T = (0, \dots, 0, 1)$ because $\mathbf{x}$ is given by

$$\mathbf{x} = x_1 \mathbf{e}_1 + \dots + x_p \mathbf{e}_p.$$

The corresponding vector of principal components is given by $\mathbf{z} = \mathbf{P}^T \mathbf{x}$, which is equivalent to $\mathbf{x} = \mathbf{P} \mathbf{z}$. Therefore the point $\mathbf{x}$ is also represented by

$$\mathbf{x} = \mathbf{P} \mathbf{z} = (\mathbf{a}_1 \dots \mathbf{a}_p) \mathbf{z} = z_1 \mathbf{a}_1 + \dots + z_p \mathbf{a}_p.$$

Therefore, $z_1, \dots, z_p$ are the coordinate values when the $\mathbb{R}^p$ is spanned by the vectors $\mathbf{a}_1, \dots, \mathbf{a}_p$. Thus the principal components describe the same point but use a different coordinate system. The system of basis vectors used by principal components $\mathbf{a}_1, \dots, \mathbf{a}_p$ is orthogonal as the commonly used system of unit vectors $\mathbf{e}_1, \dots, \mathbf{e}_p$.

The multivariate normal distribution with mean zero has the density

$$f(\mathbf{x}) = \frac{1}{(2\pi)^{p/2} |\boldsymbol{\Sigma}|^{1/2}} \exp \left\{ -\frac{1}{2} \mathbf{x}^T \boldsymbol{\Sigma}^{-1} \mathbf{x} \right\}.$$

Let us consider the spectral decomposition of $\boldsymbol{\Sigma}$ given by $\boldsymbol{\Sigma} = \mathbf{P} \boldsymbol{\Lambda} \mathbf{P}^T$. Since $\mathbf{P}$ is orthogonal the inverse of $\boldsymbol{\Sigma}$ is given by $\boldsymbol{\Sigma}^{-1} = \mathbf{P} \boldsymbol{\Lambda}^{-1} \mathbf{P}^T$. Therefore the relevant part of the density is

$$\mathbf{x}^T \boldsymbol{\Sigma}^{-1} \mathbf{x} = \mathbf{x}^T \mathbf{P} \boldsymbol{\Lambda}^{-1} \mathbf{P}^T \mathbf{x} = \mathbf{z}^T \boldsymbol{\Lambda}^{-1} \mathbf{z} = \sum_{i=1}^p \left(\frac{z_i}{\sqrt{\lambda_i}} \right)^2.$$

That means all points that have the same value of the density, that is, $\mathbf{x}^T \boldsymbol{\Sigma}^{-1} \mathbf{x} = c$ for some fixed value c can also be described as the points that fulfill

$$\sum_{i=1}^p \left(\frac{z_i}{\sqrt{\lambda_i}} \right)^2 = c.$$

These points describe in the coordinates $z_1, \dots, z_p$ an ellipsoid with lengths $\sqrt{\lambda_i c}$ on the axes.

1.2 Principal Components for Observations

For observations the problem has to be slightly modified. Let $\mathbf{x}_1, \dots, \mathbf{x}_n$ be vector valued observations, $\mathbf{x}_i^T = (x_{i1}, \dots, x_{ip})$, a linear transformation of the observations $\mathbf{x}_1, \dots, \mathbf{x}_n$ by use of vector $\boldsymbol{\alpha}_j$ yields the observations

$$z_{ij} = \boldsymbol{\alpha}_j^T \mathbf{x}_i, i = 1, \dots, n.$$

The empirical variance of the observations $z_{1j}, \dots, z_{nj}$ is given by

$$s_j^2 = \frac{1}{n} \sum_j (z_{ij} - \bar{z}_j)^2 = \boldsymbol{\alpha}_j^T \mathbf{S}_x \boldsymbol{\alpha}_j,$$

where $\bar{z}_j = \frac{1}{n} \sum_{i=1}^n z_{ij}$ and $\mathbf{S}_x = \frac{1}{n} \sum_{i=1}^n (\mathbf{x}_i - \bar{\mathbf{x}})(\mathbf{x}_i - \bar{\mathbf{x}})^T$ is the empirical covariance matrix computed from the data $\mathbf{x}_1, \dots, \mathbf{x}_n$.

The principal components for observations are then obtained by finding vectors $\boldsymbol{\alpha}_1, \dots, \boldsymbol{\alpha}_r$ such that

$$s_j^2 = \boldsymbol{\alpha}_j^T \mathbf{S}_x \boldsymbol{\alpha}_j \rightarrow \max_{\boldsymbol{\alpha}_j} \quad (1.2)$$

under the constraints $\|\boldsymbol{\alpha}_j\| = 1, \boldsymbol{\alpha}_j^T \boldsymbol{\alpha}_s = 0, s = 1, \dots, j-1$. Thus the only difference between (1.1) and (1.2) is that the covariance matrix $\boldsymbol{\Sigma}$ is replaced by the empirical covariance matrix $\mathbf{S}_x$.

Principal Components by Maximization of Empirical Variance

Find weights $\boldsymbol{\alpha}_1, \dots, \boldsymbol{\alpha}_r$ such that for $z_j = \boldsymbol{\alpha}_j^T \mathbf{x}_i$

$$\text{var}(z_j) = \boldsymbol{\alpha}_j^T \mathbf{S}_x \boldsymbol{\alpha}_j \rightarrow \max_{\boldsymbol{\alpha}_j}$$

with side constraints $\|\boldsymbol{\alpha}_j\| = 1, \boldsymbol{\alpha}_j^T \boldsymbol{\alpha}_s = 0, s = 1, \dots, j$.

The problem of maximizing the empirical covariance S is formally the same as maximizing the covariance Σ . Therefore the solution is given by the spectral decomposition of S

$$\mathbf{S} = \mathbf{Q} \mathbf{L} \mathbf{Q}^T$$

where the columns of $\mathbf{Q} = (\mathbf{q}_1, \dots, \mathbf{q}_p)$ are the eigenvectors of S and $\mathbf{L} = \text{diag}(l_1, \dots, l_p)$ is a diagonal matrix with the eigenvalues $l_1 \geq \dots \geq l_p$ of S . By use of $\mathbf{q}_1, \dots, \mathbf{q}_p$ the original data $\mathbf{x}_1, \dots, \mathbf{x}_n$ are transformed to the vector of principal components

$$\mathbf{z}_i = \mathbf{Q}^T \mathbf{x}_i.$$

It is easy to show that for the empirical covariance matrix for data $\mathbf{z}_1, \dots, \mathbf{z}_n$ one obtains

$$\mathbf{S}_z = \frac{1}{n} \sum_{i=1}^n (\mathbf{z}_i - \bar{\mathbf{z}})(\mathbf{z}_i - \bar{\mathbf{z}})^T = \mathbf{L}.$$

Thus the principal components are uncorrelated and the empirical variance of the j th principal component is given by $s_j^2 = l_j$. Moreover, in analogy to the decomposition of the underlying covariance matrix $\boldsymbol{\Sigma}$ one obtains for the empirical covariance matrix:

$$\text{tr}(\mathbf{S}_x) = \text{tr}(\mathbf{S}_z), \quad |\mathbf{S}_x| = |\mathbf{S}_z|.$$

Principal Components for Observations

The vector valued principal components are given by $z_i = Q^T x_i$, $i = 1, \dots, n$, where

$$S_x = QLQ^T$$

is the spectral decomposition with $Q = (q_1 \dots q_p)$ containing the eigenvectors of S_x as columns and $L = \text{diag}(l_1, \dots, l_p)$ being the diagonal matrix of eigenvalues of S_x .

$$S_z = \frac{1}{n-1} \sum_{i=1}^n (z_i - \bar{z})(z_i - \bar{z})^T = L$$

$$\text{tr}(S_x = \text{tr}(S_z)), |S_x| = |S_z|$$

1.3 Estimation

Principal components in the random variable and the observation case are based on the spectral decompositions

$$\Sigma = P\Lambda P^T \quad \text{and} \quad S_x = QLQ^T.$$

The corresponding eigenvectors and eigenvalues q_i, l_i from Q and L may be considered as estimates of p_i, λ_i from P and Λ . If one assumes that observations $x_1, \dots, x_n$ are iid and normally distributed one obtains for nS_x

$$nS_x \sim W(\Sigma, n-1).$$

If for the underlying eigenvalues $\lambda_1 > \dots > \lambda_p$ holds, it can be shown that asymptotically ($n \rightarrow \infty$) for the vector $\hat{\lambda}^T = (l_1, \dots, l_p)$ one has

$$\sqrt{n}(\hat{\lambda} - \lambda) \sim N(0, 2\Lambda^2).$$

For $\hat{\alpha}_j = q_j$ one obtains

$$\sqrt{n}(\hat{\alpha}_j - \alpha_j) \sim N(0, \lambda_j \sum_{\substack{s=1 \\ s \neq j}}^p \frac{\lambda_s}{(\lambda_j - \lambda_s)^2} \alpha_s \alpha_s^T)$$

(see Anderson, 2003, Section 13.5).