

Vektoren und Matrizen

**Elemente der Linearen Algebra für Multivariate Statistische
Verfahren**

U. Mortensen

Version 06. 01. 2018

Inhaltsverzeichnis

1	Einführung	6
1.1	Vektoren und Matrizen: ein intuitiver Zugang	6
1.2	Eine sehr kurze Geschichte der Vektor- und Matrixrechnung	8
2	Vektoren	10
2.1	Der euklidische Raum $\mathbb{R}^n$	10
2.2	Definition von Vektoren	12
2.2.1	Typen von Vektoren	15
2.2.2	Linearkombinationen	17
2.2.3	Produkte von Vektoren	19
2.2.4	Die Cauchy-Schwarzsche Ungleichung	28
2.2.5	Das dyadische Produkt $\mathbf{xy}'$	29
2.3	Lineare Unabhängigkeit von Vektoren	29
2.3.1	Definition der linearen Unabhängigkeit	29
2.3.2	Lineare Unabhängigkeit und Skalarprodukt	33
2.3.3	Lineare Unabhängigkeit und Korrelationen	35
2.3.4	Lineare Unabhängigkeit und Orthogonalität	36
2.4	Vektorräume	37
2.4.1	Der Begriff des Vektorraums	37
2.4.2	Basen von Vektorräumen und Teilvektorräumen	46
2.4.3	Zusammenfassung für den Fall $V = \mathbb{R}^n$	58
2.4.4	Polynome, stetige Funktionen und Vektorräume*	59
3	Matrizen	61
3.1	Definitionen	61
3.2	Elementare Operationen mit Matrizen	64
3.3	Die Multiplikation von Matrizen	65
3.3.1	Die Multiplikation einer Matrix mit einem Vektor	65
3.3.2	Der allgemeine Fall	67
3.3.3	Transformationen und Abbildungen	69
3.4	Anwendung: Mittelwerte und Varianzen	74
3.5	Matrizen und Vektorräume	76

3.6	Der Rang einer Matrix	78
3.7	Determinanten	83
3.7.1	Die Definition der Determinante	83
3.7.2	Eigenschaften der Determinante	86
3.8	Die Inverse einer $(n \times n)$ -Matrix	89
3.9	Quadratische Formen und Eigenvektoren symmetrischer Matrizen .	95
3.9.1	Rotationen	95
3.9.2	Quadratische Formen und Eigenvektoren	97
3.9.3	Das charakteristische Polynom und Eigenräume	108
3.9.4	Spektraldarstellung einer symmetrischen Matrix M	110
3.9.5	Kovarianz und generalisierte Varianz	111
3.9.6	Die Inverse einer symmetrischen Matrix	113
3.9.7	Die Wurzel aus einer positiv semidefiniten Matrix	114
3.9.8	Die Singularwertzerlegung (SVD)	115
3.10	Maximalprinzipien	119
3.10.1	Die Differentiation von Vektoren	119
3.10.2	Die Differentiation von quadratischen Formen	120
3.10.3	Die Methode der Kleinsten Quadrate	121
3.10.4	Generalisierte Kleinste Quadrate	124
3.10.5	Extrema unter Nebenbedingungen	125
3.10.6	Der Rayleigh-Quotient und seine Maximierung	127
3.10.7	Vektor- und Matrixnormen	130
3.10.8	Die Approximation von Matrizen	133
3.11	Basen und Transformationen von Basen	137
3.12	Bestimmung einer Basis für eine Datenmatrix	139
3.13	Singularwertzerlegung und PCA	142
3.14	Eigenvektorberechnung und Deflation einer Matrix	146
3.15	Die verallgemeinerte Inverse	148
3.16	Lineare Gleichungssysteme	149
3.16.1	Allgemeine Charakterisierung der Lösungen	149
3.16.2	Die Cramersche Regel	152
3.16.3	Lineare Gleichungen und Gauß-Algorithmus	153
3.16.4	Die Cholesky-Zerlegung	155

3.17	Projektionen	157
3.17.1	Orthogonale Projektion eines Vektors auf einen anderen . .	157
3.17.2	Projektionen auf Hauptachsen	158
3.17.3	Projektionen auf k -dimensionale Teilräume	158
3.17.4	Projektion eines Datenvektors auf einen Teilraum	161
3.18	Schlecht konditionierte Matrizen und Regularisierung*	163
3.19	Kroneckerprodukte	165
4	Abbildungen und Funktionen	166
4.1	Allgemeine Definition von Abbildungen	166
4.2	Lineare Abbildungen	169
4.3	Kern und Bild einer linearen Abbildung	172
4.4	Die Matrix einer linearen Abbildung	175
5	Eigenvektoren und Eigenwerte nichtsymmetrischer Matrizen	178
5.1	Der allgemeine Fall	178
5.2	Das generalisierte Eigenvektorproblem	184
5.3	Mehrfache Eigenwerte	187
6	Funktionenräume	188
6.1	Einführung	188
6.2	Normierte Räume	189
6.2.1	Definitionen	189
6.2.2	Anmerkungen zur Konvergenz in Funktionenräumen	193
6.3	Hilberträume	195
6.3.1	Definition und Eigenschaften	195
6.3.2	Lineare Operatoren	202
6.3.3	Kernfunktionen	206
7	Kernmethoden	213
7.1	Klassifikation	213
7.2	Kern-Regression	217
7.3	Kernel-PCA	218
8	Anhang	219

8.1	Das vektorielle Produkt	219
8.2	Ungleichungen	223
8.2.1	Eine allgemeine Ungleichung	223
8.2.2	Die Höldersche Ungleichung	224
8.3	Der allgemeine Begriff des Vektorraums	225
8.4	Einfache lineare Regression	228
8.5	Alternativer Beweis von Satz 3.27	230
8.6	Gleichungssysteme und Singularwertzerlegung (II)	232
8.7	Alternative Herleitung der PCA	235
8.8	Ein Maximum-Prinzip	236
8.9	Die n -dimensionale Normalverteilung	237
Literatur		238
Index		240

1 Einführung

1.1 Vektoren und Matrizen: ein intuitiver Zugang

Die Vektor- und Matrixrechnung ist ein nahezu unerlässliches Hilfsmittel zur Behandlung der Fragestellungen der multivariaten Statistik. Ohne sie lassen sich die Abhängigkeiten zwischen den untersuchten Variablen kaum in übersichtlicher Weise analysieren. Am Beispiel der multiplen Regression läßt sich dieser Sachverhalt illustrieren. Die multiple Regression ist durch die Beziehung

$$Y = b_0 + b_1 X_1 + b_2 X_2 + \cdots + b_p X_p + e \quad (1.1)$$

definiert. Dabei stehen $Y, X_1, \dots, X_p$ für Messungen der Variablen $V_y, V_1, \dots, V_p$, und die $b_0, b_1, \dots, b_p$ sind zunächst unbekannte Regressionsgewichte, deren Werte aus den Messungen geschätzt werden müssen; e repräsentiert einen "Fehler", d.h. den Effekt aller in der Untersuchung nicht weiter kontrollierten Variablen, die außer den $X_1, \dots, X_p$ noch einen Einfluß auf Y haben. Für Y und die *Prädiktoren* $X_1, \dots, X_p$ liegen m Messungen vor, so dass (1.1) einem System von m Gleichungen mit den $p + 1$ Unbekannten $b_0, b_1, \dots, b_p$ ist:

$$Y_i = b_0 + b_1 x_{i1} + b_2 x_{i2} + \cdots + b_p x_{ip} + e_i, \quad i = 1, \dots, m \quad (1.2)$$

Ein erstes Problem ist die Schätzung der Parameter $b_0, b_1, \dots, b_p$; sie wird üblicherweise mit der Methode der Kleinsten Quadrate vorgenommen. Schon mit relativ kleinen Werten von m und p werden die Rechnungen sehr schnell unübersichtlich, so dass es sich lohnt, eine kompaktere Schreibweise einzuführen. Dazu bietet sich der Gebrauch der Vektor- und Matrixschreibweise an und die dazu korrespondierenden Vektor- und Matrixrechnungen. Schreibt man nämlich das System (1.2) aus, so erhält man die Darstellung

$$\begin{aligned} Y_1 &= b_0 + b_1 X_{11} + b_2 X_{12} + \cdots + b_p X_{1p} + e_1 \\ Y_2 &= b_0 + b_1 X_{21} + b_2 X_{22} + \cdots + b_p X_{2p} + e_2 \\ &\vdots \\ Y_m &= b_0 + b_1 X_{m1} + b_2 X_{m2} + \cdots + b_p X_{mp} + e_m \end{aligned} \quad (1.3)$$

Der Parameter b_1 taucht in allen Gleichungen als Koeffizient der Messungen $X_{11}, \dots, X_{m1}$ der ersten Prädiktorvariablen auf. Man fasst nun diese Messungen zusammen, indem man sie als eine in Klammern gesetzte Spalte schreibt. In derselben Weise geht man mit der zweiten bis zur p -ten Spalte vor. Für b_0 führt man eine in Klammern gesetzte Spalte ein, die nur Einsen enthält, und die $Y_1, \dots, Y_m$ und die $e_1, \dots, e_m$ schreibt man ebenfalls als eine in Klammern gesetzte Spalte. Auf diese Weise erhält man eine Gleichung der Form

$$\begin{pmatrix} Y_1 \\ Y_2 \\ \vdots \\ Y_m \end{pmatrix} = b_0 \begin{pmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{pmatrix} + b_1 \begin{pmatrix} x_{11} \\ x_{21} \\ \vdots \\ x_{m1} \end{pmatrix} + \cdots + b_p \begin{pmatrix} x_{1p} \\ x_{2p} \\ \vdots \\ x_{mp} \end{pmatrix} + \begin{pmatrix} e_1 \\ e_2 \\ \vdots \\ e_m \end{pmatrix} \quad (1.4)$$

Der Vergleich dieser Gleichung mit den Gleichungen (1.3) macht klar, wie diese Gleichung zu verstehen ist: der Koeffizient b_j , $j = 0, 1, \dots, p$ soll mit jeder der Zahlen in der in Klammern stehenden Spalte von Zahlen multipliziert werden, und die Spalten sollen komponentenweise addiert werden, d.h. die Zahl an der i -ten Stelle einer Spalte soll zu den Zahlen an der jeweils i -ten Position der übrigen Spalten addiert werden. Führt man noch die Bezeichnungen

$$\mathbf{y} = \begin{pmatrix} Y_1 \\ Y_2 \\ \vdots \\ y_m \end{pmatrix}, \quad \vec{1} = \begin{pmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{pmatrix}, \quad \mathbf{x}_j = \begin{pmatrix} x_{1j} \\ x_{2j} \\ \vdots \\ x_{mj} \end{pmatrix}, \quad \mathbf{e} = \begin{pmatrix} e_1 \\ e_2 \\ \vdots \\ e_m \end{pmatrix}, \quad j = 1, \dots, p \quad (1.5)$$

ein, so kann (1.4) in der Form

$$\mathbf{y} = b_0 \vec{1} + b_1 \mathbf{x}_1 + b_2 \mathbf{x}_2 + \dots + b_p \mathbf{x}_p + \mathbf{e} \quad (1.6)$$

geschrieben werden. Statt der fett geschriebenen Buchstaben ist auch die Schreibweise $\vec{y}$, $\vec{x}_1, \dots, \vec{x}_p$, $\vec{e}$ üblich; hier wird hauptsächlich von der Fettschrift Gebrauch gemacht. $\vec{1}$, $\mathbf{y}$, $\mathbf{x}_1$ etc stehen also für die in Klammern stehenden Spalten von Messwerten. Diese Spalten heißen *m-dimensionale Vektoren*¹. Die Zahlen in einer Spalte sind die *Komponenten* des Vektors, und die Redeweise vom *m*-dimensionalen Vektor heißt zunächst nur, dass die jeweilige Zahlenkolumne eben *m* Elemente enthält. Sicherlich sollte man aber auch die Reihenfolge, in der die Komponenten zwischen den Klammern erscheinen, nicht verändern – das hieße ja, die Messwerte für verschiedene Objekte bzw. Personen zu vertauschen. Nicht nur die Komponenten für sich genommen, sondern auch ihre Reihenfolge definiert demnach einen Vektor.

Für eine Reihe von Fragestellungen erweist es sich als nützlich, die Vektorgleichung (1.6) noch kompakter zu schreiben, indem man zur Matrixschreibweise übergeht. Dazu fasst man die Vektoren $\vec{1}$, $\mathbf{x}_1, \dots, \mathbf{x}_p$ spaltenweise zu einer Matrix

$$X = \begin{pmatrix} 1 & x_{11} & x_{12} & \cdots & x_{1p} \\ 1 & x_{21} & x_{22} & \cdots & x_{2p} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_{m1} & x_{m2} & \cdots & x_{mp} \end{pmatrix}. \quad (1.7)$$

zusammen. Fasst man darüber hinaus die Koeffizienten $b_0, b_1, \dots, b_p$ ebenfalls als Komponenten eines Vektors $\mathbf{b}$ auf:

$$\mathbf{b} = \begin{pmatrix} b_0 \\ b_1 \\ b_2 \\ \vdots \\ b_p \end{pmatrix}, \quad (1.8)$$

¹Vektor von lat. vehi = fahren, fahren lassen reiten; vectum → Träger, Fahrer

so kann man die Vektorgleichung (1.6) als Matrixgleichung anschreiben:

$$\mathbf{y} = X\mathbf{b} + \mathbf{e}, \quad (1.9)$$

wobei der Ausdruck $X\mathbf{b}$, d.h. das Produkt der Matrix X mit dem Vektor $\mathbf{b}$, so definiert wird, dass ausgeschrieben gerade die Gleichung

$$X\mathbf{b} = b_0\vec{1} + b_1\mathbf{x}_1 + b_2\mathbf{x}_2 + \cdots + b_p\mathbf{x}_p \quad (1.10)$$

besteht: $X\mathbf{b}$ ist einfach eine kurze Schreibweise für $b_0\vec{1} + b_1\mathbf{x}_1 + b_2\mathbf{x}_2 + \cdots + b_p\mathbf{x}_p$. Es erweist sich allerdings als nützlich, diese Definition noch ausführlicher zu formulieren, was in Abschnitt 3.3.1 auch geschehen wird.

Offenbar kann man mit Vektoren rechnen: man kann einen Vektor mit einer Zahl multiplizieren, denn $b_j\mathbf{x}_j$ bedeutet ja, dass jede Komponente von $\mathbf{x}_j$ mit b_j multipliziert werden soll. Weiter kann man Vektoren addieren, indem man ihre korrespondierenden Komponenten addiert. Es gibt auch Möglichkeiten, Produkte von Vektoren zu bilden, worauf in den folgenden Abschnitten näher eingegangen wird. Ebenso kann man offenbar auch mit Matrizen rechnen: die Gleichung (1.9) legt diese Vermutung nahe. Die weitere Entwicklung der Vektor- und Matrixrechnung wird zeigen, dass man mit ihr mehr als nur abgekürzte Schreibweisen erhält.

1.2 Eine sehr kurze Geschichte der Vektor- und Matrixrechnung

Die Vektor- und Matrixrechnung ist ein Teil der linearen Algebra, deren Geschichte weit zurückreicht: vor 4000 Jahren konnte man in Babylon Gleichungen mit zwei Unbekannten lösen. 200 vChr erschien in China ein Buch – Neun Kapitel über die Kunst der Mathematik –, in dem ein allgemeiner Lösungsansatz für drei Gleichungen mit drei Unbekannten vorgestellt wurde. Die Begriffe Vektor und Matrix treten hier noch nicht auf, aber bei diesen frühen Arbeiten wird gewissermaßen die Basis für die spätere Entwicklung dieser Begriffe gelegt. Es war dann Gottfried Wilhelm Leibniz (1646 – 1716), der im Zusammenhang mit der Beantwortung von Fragen zur Entwässerung von Gruben im Harz Gleichungssysteme mit mehr als drei Unbekannten lösen mußte, dafür eine Systematik suchte und in diesem Zusammenhang den Begriff der Determinante entwickelte, der heute im Rahmen der Matrixrechnung eingeführt wird; wohl unabhängig von Leibniz entwickelte 1683 der japanische Mathematiker Takakazu Shinzuke Seki (1642 – 1708) ebenfalls den Begriff der Determinante. Der Schweizer Mathematiker Gabriel Cramer (1704 – 1752) führte dann ca 50 Jahre später auf der Basis der Leibnizschen Vorarbeiten die nach ihm benannte Cramersche Regel ein, die eine systematische Lösung von linearen Gleichungssystemen mit mehreren Unbekannten gestattet, falls dieses System bestimmte Eigenschaften besitzt. Der ebenfalls schweizerische Mathematiker Leonhard Euler (1707 – 1783) zeigte dann, dass die Gleichungssysteme nicht notwendig eine Lösung haben. Gegen Ende des 18-ten

Jahrhunderts fand Carl Friedrich Gauß (1777 – 1855) auf der Basis der Vorarbeiten von Leibniz, Cramer und Euler eine effektive Methode zur Lösung linearer Gleichungen. Auf die Theorie linearer Gleichungssysteme wird in Abschnitt 3.16 eingegangen, und Determinanten werden kurz in Abschnitt 3.7 besprochen, – in diesem Skript wird mehr auf den Begriff der linearen Abhängigkeit von Vektoren fokussiert, da er für die Anwendung der Vektor- und Matrixrechnung in der multivariaten Statistik eine zentrale Rolle spielt.

Um die Determinanten genauer diskutieren zu können, mußte ein gewisser formaler Apparat geschaffen werden. Ein entsprechender Ansatz wurde im Jahr 1850 von dem britischen Mathematiker James Joseph Sylvester (1814 – 1897) gemacht, der den Ausdruck 'Matrix' (lat. für Gebärmutter, Mutterleib) für Felder von Zahlen der Form (1.7) einführte. Die wichtige Operation der Multiplikation von Matrizen wurde 1855 von dem britischen Mathematiker Arthur Cayley (1821 – 1895) im Jahr 1855 definiert.

Im 19-ten Jahrhundert stand aber der Begriff des Vektors wegen seiner Anwendung in der Physik mehr im Zentrum mathematischer Entwicklungen als der der Matrix. Zunächst wurde der Vektorbegriff von dem irischen Mathematiker und Physiker William Rowan Hamilton (1805 – 1865) im Zusammenhang mit der Repräsentation komplexer Zahlen $z = x + iy$, $x, y \in \mathbb{R}$, $i = \sqrt{-1}$, eingeführt; auf ihn geht auch die Bezeichnung *Skalar* als einer einzelnen reellen Zahl zurück: sie ist in Wert auf der reellen Skala zwischen $-\infty$ bis $+\infty$. C. F. Gauß, W. R. Hamilton und der deutsche Mathematiker Hermann Grassmann (1809–1877) verallgemeinerten um 1844 den Vektorbegriff für den Fall von $n > 3$ Dimensionen. Der amerikanische Physiker Josiah Willard Gibbs (1839–1903) entwickelte wesentlich die Vektoranalysis weiter, bei der die Vektorkomponenten Funktionen der Zeit sind, worauf allerdings in diesem Skript nicht eingegangen wird. Weiteren Aufschwung erfuhr der Vektor- und Matrixkalkül nach dem zweiten Weltkrieg, als größere Computer zur Verfügung standen, mit denen sich ausgedehntere Rechnungen durchführen lassen. Mittlerweile ist die Vektor- und Matrixrechnung für die multivariate Statistik ein nahezu unentbehrliches Hilfsmittel geworden.

Der Ausdruck *Matrixkalkül* drückt einen wesentlichen Aspekt der Matrixrechnung aus: die Regeln der Verknüpfung von Matrizen und Vektoren sind so angelegt, dass ihre formale Anwendung zu Einsichten in die Struktur von Transformationen z.B. von Vektoren und damit von Beziehungen zwischen Variablen die sich ohne diesen Kalkül nur sehr mühselig entschlüsseln lassen. Das Kalkülhafte der Matrix- und Vektorrechnung erleichtert diese Analysen ganz außerordentlich und führt gleichzeitig zu einem tieferen Verständnis der Beziehungen zwischen den verschiedenen Transformationen, was wiederum ein besseres Verständnis der multivariaten Verfahren bedeutet.

2 Vektoren

”Der Punkt ist für den Empirismus eine problematische Sache.”²

2.1 Der euklidische Raum $\mathbb{R}^n$

Der in Abschnitt 1.1 gegebenen intuitiven Einführung zufolge sind Vektoren sogenannte n -Tupel reeller Zahlen, bei denen es auf die Reihenfolge der Zahlen ankommt, denn würde man die Reihenfolge vertauschen, würde man den Personen oder allgemein den Objekten falsche Maßzahlen zuordnen.

Bevor aber die allgemeine Definition des n -dimensionalen Vektors vorgestellt wird, soll der n -dimensionale Punktraum definiert werden. Auf einer Geraden kann zunächst ein Nullpunkt festgelegt werden. Die Punkte rechts von diesem Punkt seien Punkte auf der ”positiven” Halbgeraden, die auf der linken Seite seien ”negative” Punkte. Der Abstand eines Punktes P auf der positiven Halbgeraden vom Nullpunkt sei $x \in \mathbb{R}$, d.h. ein Punkt entspricht einer reellen Zahl, und $\mathbb{R}$ ist die Vereinigung der Menge der ganzen Zahlen $\dots -3, -2, 1-, 0, 1, 2, 3, \dots$, der rationalen Zahlen p/q und $-p/q$, wobei $p \in \mathbb{N}, 0 \neq q \in \mathbb{N}$, und der Menge der irrationalen (= nicht als Quotient, d.h. Ratio, darstellbaren) Zahlen wie $\sqrt{2}, \pi, e$ etc. Man kann diese Gleichsetzung von Punkt und reeller Zahl als Definition des Begriffs ’Punkt’ betrachten; die philosophischen und mathematischen Betrachtungen zum Begriff des Punktes haben eine über 2000-jährige Geschichte, auf deren Details hier weder eingegangen werden kann noch muß (vergl. Bedürftig et al. (2012)). x ist dann eine Koordinate des Punktes auf der Geraden. Analog dazu ist $-x$ die Koordinaten eines Punktes auf der linken Halbgeraden. Man kann $P(x)$ oder $P = (x)$ schreiben, um anzugeben, dass p die Koordinate x auf der Geraden hat.

Gegeben sei eine Ebene, auf der ein Koordinatensystem eingetragen worden sei. Irgendein Punkt P auf der Ebene kann dann durch die Koordinaten (x_1, x_2) beschrieben werden; x_1 ist die Koordinate auf der ersten Koordinatenachse, x_2 ist die Koordinate auf der zweiten Koordinatenachse. Man kann $P(x_1, x_2)$ oder $P = (x_1, x_2)$ schreiben, um anzugeben, dass P die Koordinaten x_1 und x_2 hat. Ein Punkt entspricht nun einem Element von $\mathbb{R} \times \mathbb{R} = \mathbb{R}^2$, und $\mathbb{R}^2$ ist die Menge aller Paare von reellen Zahlen.

Analog dazu schreibt man $P(x_1, x_2, x_3)$ oder $P = (x_1, x_2, x_3)$, um den Ort eines Punktes im dreidimensionalen Raum zu spezifizieren. Man kann diese Definition auf den allgemeinen Fall von n Koordinaten verallgemeinern: $P(x_1, \dots, x_n)$ charakterisiert einen ”Punkt” im n -dimensionalen Punktraum $\mathbb{R}^n$, womit das n -fache Cartesische Produkt

$$\mathbb{R} \times \dots \times \mathbb{R} = \mathbb{R}^n. \quad (2.1)$$

²Bedürftig & Murawski (2012), p. 175

gemeint ist. Dies bedeutet, dass der n -dimensionale Punktraum die Menge aller n -Tupel $x_1 \in \mathbb{R}, x_2 \in \mathbb{R}, \dots, x_n \in \mathbb{R}$ sein soll; die x_j können alle Werte aus $\mathbb{R}$ annehmen. Dass man für $n > 3$ keine geometrische Anschauung mehr hat, ist dabei unwesentlich, da sich alle Formeln, z.B. für die Distanz zwischen irgendzwei Punkten, für beliebiges $n \in \mathbb{N}$ anschreiben lassen. Für die euklidische Distanz zwischen zwei Punkten $P(x, \dots, x_n)$ und $Q(y_1, \dots, y_n)$ gilt zum Beispiel

$$d(P, Q) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2}; \quad (2.2)$$

dies ist der Satz des Pythagoras für n Dimensionen. Es ist kein Problem, statt der Menge $\mathbb{R}$ die Menge $\mathbb{C}$ der komplexen Zahlen $z = x + iy$ mit $x, y \in \mathbb{R}, i = \sqrt{-1}$ anzunehmen, aber darauf muß in den Standardanwendungen der Vektorrechnung in der multivariaten Statistik nicht weiter eingegangen werden.

In der Einführung sind Vektoren ebenfalls als n -Tupel reeller Zahlen eingeführt worden, so dass es den Anschein haben kann, dass Punkte und Vektoren ein- und dasselbe zu sein scheinen. Dies ist nicht ganz so, wie in Abschnitt 2.2 deutlich werden wird.

Mit den Koordinaten von Punkten lassen sich gewisse Operationen durchführen, die wiederum Punkte definieren:

Definition 2.1 *Es sei $P(x_1, \dots, x_n)$ ein Punkt mit den Koordinaten $x_1, \dots, x_n$. Dann ist $Q = aP$, $a \in \mathbb{R}$ ein Punkt mit den Koordinaten*

$$Q = aP(x_1, \dots, x_n) = (ax_1, \dots, ax_n). \quad (2.3)$$

Definition 2.2 *Offenbar ist Q ein Punkt auf der Geraden, die durch den Nullpunkt O und durch P geht. Liegen mehrere Punkte auf einer Geraden, so heißen sie kollinear.*

So seien $A(x_1, y_1), B(x_2, y_2), C(x_3, y_3)$ irgend drei Punkte. Die Punkte sind kollinear, wenn

$$C - A = a(B - A), \quad a \in \mathbb{R},$$

d.h.

$$(x_3 - x_1, y_3 - y_1) = a(x_2 - x_1, y_2 - y_1).$$

Definition 2.3 *Es seien P und Q mit den Koordinaten $(x_1, \dots, x_n)$ und $(y_1, \dots, y_n)$. Dann ist $R = P + Q$ der Punkt mit den Koordinaten*

$$R = P + Q = (x_1 + y_1, \dots, x_n + y_n). \quad (2.4)$$

Die folgenden Aussagen ergeben sich direkt aus den Rechenregeln für reelle Zahlen:

1. $P + Q = Q + P$ für beliebige Punkte P und Q aus $\mathbb{R}^n$,
2. $(P + Q) + R = P + (Q + R)$
3. O_n sei der *neutrale Punkt* des $\mathbb{R}^n$, wenn $P + O_n = P$ für $P \in \mathbb{R}^n$,
4. Für irgendeinen Punkt P gilt dann $P + (-P) = P - P = O_n$,
5. $1P = P$, $P \in \mathbb{R}^n$,
6. $(ab)P = a(bP)$, $P \in \mathbb{R}^n$, $a, b \in \mathbb{R}$,
7. $(a + b)P = aP + bP$, $a, b \in \mathbb{R}$, $P \in \mathbb{R}^n$,
8. $a(P + Q) = aP + aQ$, $a \in \mathbb{R}$, $P, Q \in \mathbb{R}^n$.

Im $\mathbb{R}^n$ kann eine *Metrik* durch eine Distanz d zwischen irgendzwei Punkten erklärt werden:

Definition 2.4 Es seien $P, Q \in \mathbb{R}^n$ irgendzwei Punkte, und d sei die Distanz zwischen P und Q : Wenn die Bedingungen

1. $d(P, Q) \geq 0$,
2. $d(P, Q) = d(Q, P)$
3. Es sei $R \in \mathbb{R}^n$ ein weiterer Punkt; dann gilt $d(P, R) \leq d(P, Q) + d(Q, R)$ (Dreiecksungleichung).

Dann definiert die Distanz d eine Metrik.

Die in (2.2) erklärte euklidische Distanz d definiert die *euklidische Metrik*. Eine Verallgemeinerung der euklidischen Metrik ist die Minowski-Metrix³

$$d(P, Q) = \left(\sum_{i=1}^n (x_i - y_i)^p \right)^{1/p}, \quad p \in \mathbb{R}, \quad (2.5)$$

die insbesondere in der Multidimensionalen Skalierung Anwendung findet; Für $p = 2$ erhält man die Euklidische Metrik.

2.2 Definition von Vektoren

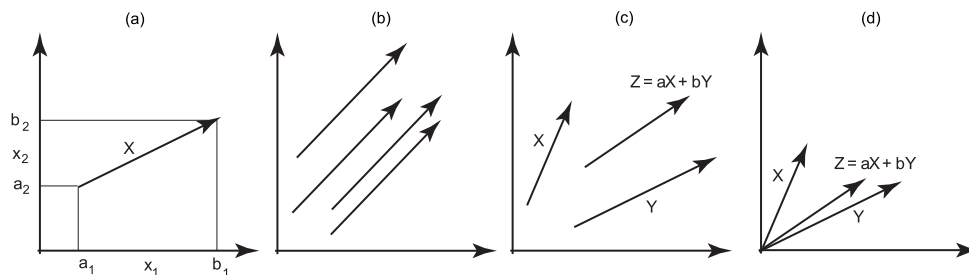
In der Einführung wurden Vektoren⁴ als geordnete n -Tupel von reellen Zahlen eingeführt; damit ergibt sich zunächst kein Unterschied zu den Punkten im $\mathbb{R}^n$, die ja ebenfalls durch n -Tupel reeller Zahlen definiert wurden. Diese Charakterisierung entspricht nicht der allgemeinsten Definition von Vektoren, aber die allgemeine Definition soll auf spätere Kapitel verschoben werden. Der Unterschied zu einem Punkt im $\mathbb{R}^n$ besteht in der Bedeutung, die dem n -Tupel zugeschrieben wird: Vektoren sind *gerichtete Größen*⁵. Dieser Ausdruck kommt aus der Physik,

³Hermann Minkowski (1864 – 1909), Mathematiker

⁴von lat. Träger, Fahrer

⁵(Wer es lieber gesungen mag: <https://www.youtube.com/watch?v=TzaYsyNvvZA>. Der Autor dieses Skripts ist nicht der Sänger.

Abbildung 1: Vektoren: (a) Komponenten, (b) verschiedene Repräsentanten eines Vektors, (c) Linearkombination, (d) der Anfangspunkt wurde in den Koordinatenursprung gelegt



wo man z.B. Kräfte betrachtet, die in jeweils eine bestimmte Richtung mit einer bestimmten Ausprägung wirken und die man durch einen Pfeil repräsentieren kann, der in die Richtung der Wirkung der Kraft zeigt und dessen Länge die Ausprägung oder Größe der Kraft repräsentiert. Ein anderes Beispiel ist ein Partikel, dass sich im $\mathbb{R}^3$ bewegt und dessen Bewegung zu einem bestimmten Zeitpunkt t durch eine bestimmte Orientierung (Richtung) und eine bestimmte Geschwindigkeit bestimmt ist. Für jeden Zeitpunkt t kann man dann einen Pfeil zeichnen, dessen Orientierung die Richtung der Bewegung anzeigt und dessen Länge die Geschwindigkeit repräsentiert. Da es nur auf die Orientierung und die Länge des Pfeils ankommt, dienen die Komponenten auch nur zur Bestimmung der Orientierung und der Länge des Pfeils, der *Ort* im Koordinatensystem wird damit nicht festgelegt. Es seien (a_1, a_2, a_3) die Koordinaten des Anfangspunktes und (b_1, b_2, b_3) seien die Koordinaten des Endpunktes des Pfeils. Die *Komponenten* des Vektors, den der Pfeil repräsentiert, sind dann

$$x_i = b_i - a_i, \quad i = 1, \dots, n \quad (2.6)$$

wobei im Falle der Beispiele aus der Physik $n = 3$. Es sei angemerkt, dass im 2-dimensionalen Fall die Orientierung θ eines Vektors durch

$$\tan \theta = \frac{x_2}{x_1} \quad (2.7)$$

gegeben ist (vergl. Abbildung 3).

Man schreibt

$$\vec{x} = \mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}. \quad (2.8)$$

Der Fall $n = 1$ ist ein Spezialfall, der als *Skalar* bezeichnet wird, weil er auf der "Skala" von $-\infty$ bis $+\infty$ liegt. Es kommt auf die Anordnung der Komponenten an:

die gleichen Komponenten in anderer Anordnung definieren einen anderen Vektor. Der Repräsentation durch einen Pfeil entspricht die in Gleichung (2.8) bereits eingeführte Schreibweise $\vec{x}$ für einen Vektor. Insbesondere im angelsächsischen Sprachraum wird oft die Fettschrift $\mathbf{x}$, $\mathbf{y}$, $\mathbf{a}$, etc gewählt. Diese Schreibweise wird hier übernommen, weil sie im Allgemeinen sehr übersichtlich ist, aber auch von der $\vec{x}$ -Schreibweise wird gelegentlich Gebrauch gemacht.

Da in die Definition eines Vektors nur die in (2.6) definierten Komponenten, nicht aber die Koordinaten a_i, b_i eingehen, kommt es offenbar auf die Position eines Vektors in einem Koordinatensystem nicht an. Dies bedeutet, dass ein Vektor eine Äquivalenzklasse definiert: ein Vektor ist die Menge aller Pfeile mit gleicher Orientierung und Länge, unabhängig von der Position eines Pfeils. Ein bestimmter Pfeil aus dieser Äquivalenzklasse ist ein *Repräsentant* des Vektors.

Beziehung zum $\mathbb{R}^n$: Da $\mathbf{x}$ ein n -Tupel ist, kann $\mathbf{x}$ ebenso einen Punkt mit den Koordinaten $(x_1, \dots, x_n)$ bedeuten. Wählt man eine Repräsentation von $\mathbf{x}$ derart, dass der Anfangspunkt des Pfeils im Ursprung des Koordinatensystems liegt, so definiert $(x_1, \dots, x_n)$ tatsächlich den Endpunkt der Repräsentation. Soll also $\mathbf{x}$ einen Vektor bezeichnen, so ist offenbar etwas anderes gemeint, nämlich nicht einen Punkt, sondern ein Linienelement mit einer bestimmten Länge und Orientierung. Diese Interpretation erleichtert bestimmte Betrachtungen. Dass man, wie in der Einführung angedeutet, Messwerte als Komponenten wählen kann, erweist sich als mit dieser Interpretation kompatibel, wie noch deutlich werden wird. $\square$

Spalten- und Zeilenvektoren: Ein Vektor wird, wie in (2.8) bereits angedeutet, als *Spalte* angeschrieben. Ein Vektor kann *gestürzt* oder *transponiert* werden, – dann wird er als Zeile angeschrieben und mit $\mathbf{x}'$ bezeichnet. Es ist dann

$$\vec{x}' = \mathbf{x}' = \mathbf{x}^T = \mathbf{x}^t = \mathbf{x}^\top = (x_1, x_2, \dots, x_n). \quad (2.9)$$

$\mathbf{x}'$ und $\mathbf{x}^T$ sind nur verschiedene Schreibweisen, T bzw. t stehen für 'transponiert'. Damit es bei der Verwendung von T oder t nicht zu Verwechslungen kommt, wird auch ein spezielles T wie in $\mathbf{x}^\top$ verwendet. In diesem Text wird hauptsächlich die Schreibweise $\mathbf{x}'$ verwendet. Natürlich ist dann $(\mathbf{x}')' = \mathbf{x}$. Dementsprechend schreibt man auch zur Platzersparnis $\mathbf{x} = (x_1, x_2, \dots, x_n)'$, d.h. ein Zeilenvektor, der gestürzt wird, ist wieder ein Spaltenvektor.

Da die Komponenten als Koordinatendifferenzen zu interpretieren sind und der Ort des Vektors keine Rolle spielt (vergl. Abbildung 1), werden in vielen Anwendungen die Anfangspunkte der Vektoren in den Koordinatenursprung (0-Punkt) gelegt. Die Komponenten entsprechen dann den Koordinaten des Endpunktes eines Vektors.

Spezielle Vektoren: Es gibt spezielle Vektoren, die sich in Anwendungen des

Vektorbegriffs als nützlich erweisen:

$$\vec{0} = (0, 0, \dots, 0)' \quad (2.10)$$

$$\vec{1} = (1, 1, \dots, 1)' \quad (2.11)$$

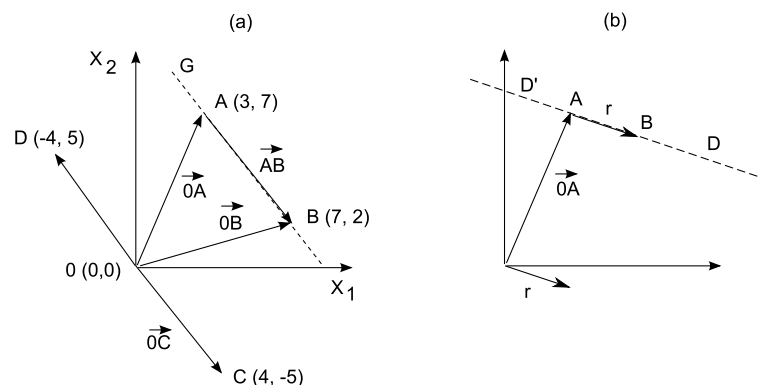
$$\mathbf{e}_i = (0, \dots, 0, 1, 0, \dots, 0)' \quad (2.12)$$

$\vec{0}$ ist der *Nullvektor*, seine Komponenten sind alle gleich Null. $\vec{1}$ ist der *Einsvektor*, seine Komponenten sind alle gleich 1. $\mathbf{e}_i$ ist der *i-te Einheitsvektor*. Seine Komponenten sind alle gleich Null bis auf die *i*-te, die gleich 1 ist. Die Anzahl der Komponenten dieser Vektoren ist gleich n , und der Wert von n wird durch den Zusammenhang, in dem diese Vektoren verwendet werden, festgelegt. Für einen gegebenen Wert von n gibt es n Einheitsvektoren, da i den Wert $1, 2, \dots, n$ annehmen kann.

2.2.1 Typen von Vektoren

Gelegentlich ist es nützlich, bestimmte Typen von Vektoren zu unterscheiden. Gegeben sei ein Koordinatensystem; in Abbildung 2 wird der Übersichtlichkeit wegen nur ein 2-dimensionales Koordinatensystem betrachtet, aber alle Begriffe übertragen sich auf n -dimensionale Systeme. In einem Koordinatensystem werden Punkte durch Angaben ihrer Koordinaten spezifiziert, etwa: $A: (x_1, \dots, x_n)$; die Koordinaten werden stets in einer Zeile angeschrieben. In Abb. 2 (a) hat man die Punkte 0 mit den Koordinaten $(0, 0)$, $A: (3, 7)$, $B: (7, 2)$, und $C: (4, -5)$.

Abbildung 2: Orts-, Verbindungs-, Richtungs- und Stützvektoren



Ortsvektoren: Man kann nun Vektoren betrachten, die vom Nullpunkt des Koordinatensystems zu einem Punkt weisen: $\vec{OA}$, $\vec{OB}$ und $\vec{OD}$. Diese Vektoren heißen *Ortsvektoren*, eben weil sie zu einem Ort, d.h. einem Punkt zeigen. Sie

werden wie üblich in Spaltenform geschrieben:

$$\overrightarrow{OA} = \begin{pmatrix} 3 \\ 7 \end{pmatrix}, \overrightarrow{OB} = \begin{pmatrix} 7 \\ 2 \end{pmatrix}, \overrightarrow{OC} = \begin{pmatrix} 4 \\ -5 \end{pmatrix}. \quad (2.13)$$

Verbindungsvektoren: Man kann auch den Vektor betrachten, dessen Anfangspunkt der Punkt A und dessen Endpunkt der Punkt B ist; Dies ist ein *Verbindungsvektor*. Er ist durch Komponenten definiert, die sich als Differenzen der Koordinaten der Punkte B und A ergeben

$$\overrightarrow{AB} = \begin{pmatrix} 7-3 \\ 2-7 \end{pmatrix} = \begin{pmatrix} 7 \\ 2 \end{pmatrix} - \begin{pmatrix} 3 \\ 7 \end{pmatrix} = \begin{pmatrix} 4 \\ -5 \end{pmatrix} = \overrightarrow{OB} - \overrightarrow{OA} \quad (2.14)$$

oder, analog dazu, der Punkte A und B

$$\overrightarrow{BA} = \begin{pmatrix} 3-7 \\ 7-2 \end{pmatrix} = \begin{pmatrix} 3 \\ 7 \end{pmatrix} - \begin{pmatrix} 7 \\ 2 \end{pmatrix} = \begin{pmatrix} -4 \\ 5 \end{pmatrix} = \overrightarrow{OA} - \overrightarrow{OB}.$$

Die Komponenten des Verbindungsvektors $\overrightarrow{AB}$ sind offenbar gleich den Komponenten des Ortsvektors für den Punkt C , der wiederum als Verbindungsvektor für die Punkte O und C angesehen werden kann; $\overrightarrow{OC}$ und $\overrightarrow{AB}$ sind verschiedene Repräsentanten desselben Vektors. Natürlich kann man auch den Verbindungsvektor vom Punkt B zum Punkt A bestimmen:

$$\overrightarrow{BA} = \begin{pmatrix} 3-7 \\ 7-2 \end{pmatrix} = \begin{pmatrix} 3 \\ 7 \end{pmatrix} - \begin{pmatrix} 7 \\ 2 \end{pmatrix} = \begin{pmatrix} -4 \\ 5 \end{pmatrix} = \overrightarrow{OA} - \overrightarrow{OB} \quad (2.15)$$

Offenbar ist $\overrightarrow{AB} = -\overrightarrow{BA}$, d.h. die beiden Vektoren zeigen in entgegengesetzte Richtungen, wie die Vektoren $\overrightarrow{OC}$ und $\overrightarrow{OD} = \overrightarrow{BA}$.

Richtungsvektoren: Alle Vektoren können auch als *Richtungsvektoren* oder Orientierungsvektoren aufgefasst werden: $\overrightarrow{AB}$ bzw. $\overrightarrow{BA}$ geben die Richtung bzw. Orientierung der Geraden G an. Allgemein zeigt jeder Vektor die Orientierung einer Geraden an, auf der der Vektor liegt. Man kann etwa den Vektor $\overrightarrow{OC}$ also Orientierungsvektor für die Gerade G auffassen; man schreibt dann auch $\vec{r} = \mathbf{r} = \overrightarrow{OC}$.

Um etwa eine Gerade darzustellen, benötigt man einen Punkt, der auf der Geraden G liegt und damit die Position der Geraden festlegt, und einen Richtungsvektor $\mathbf{r}$, der die Orientierung von G festlegt: mit $\mathbf{x}_A = \overrightarrow{OA}$ hat man

$$G : \mathbf{x}_A + \lambda \mathbf{r}, \quad \lambda \in \mathbb{R} \quad (2.16)$$

wobei λ ein Parameter ist, mit dem ein Punkt D oder D' auf der Geraden bestimmt wird. In Bezug auf Abbildung 2 (a) hat man insbesondere für den Punkt B

$$G : \begin{pmatrix} 3 \\ 7 \end{pmatrix} + 1 \cdot \begin{pmatrix} 4 \\ -5 \end{pmatrix} = \begin{pmatrix} 7 \\ 2 \end{pmatrix}, \quad \lambda = 1.$$

Die Beziehung (2.16) definiert also Ortsvektoren für Punkte, die auf der Geraden liegen. Der Richtungsvektor $\mathbf{r}$ kann natürlich beliebig vorgegeben werden; hier ist er nur speziell für eine Orientierung gewählt worden, die durch die Punkte A und B festgelegt wurde. Abbildung 2 (b) illustriert den Fall, bei dem ein Punkt A auf der Geraden durch den Ortsvektor $\overrightarrow{OA}$ definiert wird und die Orientierung der Geraden durch den Richtungsvektor $\mathbf{r} = (3, -1)'$ (man beachte, dass $\mathbf{r}$ hier als Zeilenvektor $(3, -1)'$ angeschrieben wird, es ist also nicht der Punkt $(3, -1)$ gemeint).

Stützvektoren: Die Ortsvektoren $\overrightarrow{OA}$ und $\overrightarrow{OB}$ – die ja auch Verbindungsvektoren für die Punkte O und A bzw. B sind – können als *Stützvektoren* für die Gerade G aufgefasst werden. Ihre Endpunkte definieren eine Gerade, nämlich die Gerade, die durch die Endpunkte dieser Vektoren verläuft, das ist hier die Gerade G , – die Vektoren stützen gewissermaßen die Gerade. Eine Gerade kann auch durch für einen Stützvektor, etwa $\overrightarrow{OA}$ definiert werden, sofern noch ein Richtungs- oder Orientierungsvektor gegeben ist, s. Abbildung 2 (b).

Im 3-dimensionalen Raum wird man drei 3-dimensionale Vektoren benötigen, um eine Ebene zu "stützen", d.h. um ihrer Orientierung im Raum festzulegen, und im n -dimensionalen Raum wird man n n -dimensionale Vektoren benötigen, um eine Verallgemeinerung der Ebene, eine *Hyperebene*, orientierungsmäßig zu spezifizieren (Hyperebenen werden in Abschnitt 2.4.1 betrachtet).

Zufallsvektoren: Die Komponenten eines Vektors können zufällige Variablen sein; der Vektor heißt dann auch *Zufallsvektor*. So ist ein Vektor, dessen Komponenten Messwerte sind, immer auch ein Zufallsvektor. Zufällige Veränderliche X haben unter bestimmten Bedingungen einen Erwartungswert $\mathbb{E}(X)$, das ist das arithmetische Mittel über alle möglichen Realisierungen von X . Dementsprechend kann ein zufälliger Vektor $\mathbf{x} = (x_1, \dots, x_n)'$ mit einem Erwartungswertvektor $\mathbb{E}(\mathbf{x}) = (\mathbb{E}(x_1), \dots, \mathbb{E}(x_n))'$ assoziiert werden. In diesem Skript wird im Allgemeinen nicht besonders hervorgehoben, dass ein Vektor ein Zufallsvektor ist, da es hier mehr auf die Algebra von Vektoren ankommt.

Zusammenfassend kann man sagen, dass jeder Vektor als Orts-, Verbindungs-, Richtungs-, Stütz- oder Zufallsvektor aufgefasst werden kann; welche Bezeichnung man wählt, hängt von der Funktion ab, die man einem Vektor bei der jeweiligen Betrachtung zuordnet.

2.2.2 Linearkombinationen

Vektoren können unter bestimmten Bedingungen addiert werden, wobei die Vektoren noch "gewichtet" werden können, wobei die Gewichtung formal der *Multiplikation eines Vektors mit einem Skalar* entspricht. Dazu sei $\lambda \in \mathbb{R}$ ein Skalar,

und $\mathbf{x}$ sei ein Vektor. Dann ist mit $\lambda\mathbf{x}$ der Vektor

$$\lambda\mathbf{x} = \lambda \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} \lambda x_1 \\ \lambda x_2 \\ \vdots \\ \lambda x_n \end{pmatrix}, \quad (2.17)$$

d.h. die Multiplikation mit einem Skalar bedeutet, dass jede Komponente des Vektors mit diesem Skalar multipliziert wird; in Gleichung (1.4) wurde diese Schreibweise bereits eingeführt. Eine geometrische Veranschaulichung des Effekts dieser Multiplikation wird in Abschnitt 2.2.3 gegeben.

Es seien $\mathbf{x}_1$ und $\mathbf{x}_2$ zwei n -dimensionale Vektoren, und es seien λ und μ irgendzwei Skalare. Dann heißt der Vektor $\mathbf{y}$, der durch

$$\mathbf{y} = \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix} = \lambda\mathbf{x}_1 + \mu\mathbf{x}_2 = \begin{pmatrix} \lambda x_{11} + \mu x_{12} \\ \lambda x_{21} + \mu x_{22} \\ \vdots \\ \lambda x_{n1} + \mu x_{n2} \end{pmatrix} \quad (2.18)$$

definiert ist, eine *Linearkombination* der Vektoren $\mathbf{x}_1$ und $\mathbf{x}_2$. Man beachte, dass hier der Begriff "linear" gerechtfertigt ist: in (2.18) tritt keine von $\mathbf{x}_1$ und $\mathbf{x}_2$ unabhängige Konstante bzw. kein Konstantenvektor auf; gäbe es einen solchen Vektor in (2.18), so wäre die Beziehung affin.

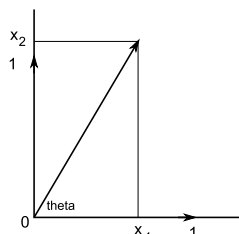
Die Gleichung (1.6), die die multiple Regression als Vektorgleichung darstellt, ist also formal eine Linearkombination. Mit dem Begriff der Linearkombination wird auch erklärt, was unter der Summe und der Differenz zweier Vektoren zu verstehen ist. Sei $\lambda = \mu = 1$. Dann ist $\mathbf{y}$ die Summe der beiden Vektoren, und die Komponenten von $\mathbf{y}$ sind die Summen der Komponenten von $\mathbf{x}_1$ und $\mathbf{x}_2$. Die Differenz erhält man, wenn man $\lambda = 1$ und $\mu = -1$ setzt.

Beispiel 2.1 Linearkombinationen von Einheitsvektoren Es seien $\mathbf{e}_i$ die in (2.12) eingeführten Einheitsvektoren, $i = 1, \dots, n$. $\mathbf{x} = (x_1, x_2, \dots, x_n)'$ sei ein beliebiger n -dimensionaler Vektor. Dann ist $\mathbf{x}$ die Linearkombination der Einheitsvektoren $\mathbf{e}_1, \dots, \mathbf{e}_n$. Denn

$$\mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} = x_1 \begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix} + x_2 \begin{pmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{pmatrix} + \dots + x_n \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{pmatrix}. \quad (2.19)$$

Abb. 3 zeigt den Vektor $\mathbf{x}$ in einem rechtwinkligen Koordinatensystem, in dem der Einheitsvektor $\mathbf{e}_1 = (1, 0)'$ auf der X_1 -Achse und der Einheitsvektor $\mathbf{e}_2 = (0, 1)'$ auf der X_2 -Achse liegt. Tatsächlich stehen die Einheitsvektoren senkrecht aufeinander; die Begründung für diese Aussage wird in Abschnitt 2.2.3, Beispiel

Abbildung 3: Vektor als Linearkombination von Einheitsvektoren; $\tan \theta = x_2/x_1$ ist die Orientierung des Vektors.



2.5, Seite 26 geliefert. Die Komponenten x_1 und x_2 von $\mathbf{x}$ erscheinen hier als Koordinaten des Endpunkts von $\mathbf{x}$ in einem durch $\mathbf{e}_1$ und $\mathbf{e}_2$ definierten Koordinatensystem. Die Interpretation der Komponenten als Koordinaten in einem durch Vektoren bestimmten Koordinatensystem wird in allgemeiner Form in Abschnitt 2.4.2 diskutiert. $\square$

2.2.3 Produkte von Vektoren

Die Regeln (i) der Multiplikation eines Vektors mit einem Skalar und (ii) der Addition von Vektoren, die wegen der Regel (i) auch die Subtraktion von Vektoren erklärt (man muß einen Vektor nur mit -1 multiplizieren und die Summation wird zu einer Subtraktion), legen nahe, auch die Multiplikation von Vektoren zu definieren. Analog zur Addition von Vektoren – die Summe zweier Vektoren wird durch die Summe der zueinander korrespondierenden Komponenten erklärt – könnte man vereinbaren, ein Produkt $\mathbf{x} * \mathbf{y}$ der Vektoren $\mathbf{x}$ und $\mathbf{y}$ durch

$$\mathbf{x} * \mathbf{y} = (x_1y_1, x_2y_2, \dots, x_ny_n)'$$

zu definieren, also durch einen Vektor, dessen Komponenten die Produkte der zueinander korrespondierenden Komponenten sind. Eine solche Definition wäre formal möglich, aber es zeigt sich, dass ihr keine interessante Bedeutung zuzuordnen ist. Es haben sich drei Definitionen von Produkten von Vektoren als interessant erwiesen:

- (i) Das *Skalarprodukt* (auch: inneres Produkt, oder dot product im Englischen) $\mathbf{x}'\mathbf{y}$, das gleich einem Skalar, also einer einzelnen reellen Zahl ist. Es spielt in der multivariaten Statistik eine zentrale Rolle, u.a. erweisen sich z. B. Mittelwerte, Varianzen und Kovarianzen als Skalarprodukte von Vektoren. Es wird im folgenden Abschnitt 2.2.3 vorgestellt.
- (ii) das *dyadische Produkt* $\mathbf{xy}'$, das gleich einer Matrix ist, deren Elemente durch Produkte der Komponenten von $\mathbf{x}$ und $\mathbf{y}$ definiert sind. Das Produkt $\mathbf{xy}'$

erweist sich als sehr nützlich, wenn Varianz-Kovarianzmatrizen definiert werden; es wird erst in Abschnitt 2.2.5 besprochen, weil es den Matrixbegriff voraussetzt.

- (iii) Das *vektorielle Produkt* (auch: Kreuzprodukt, oder äußeres Produkt) $\mathbf{x} \times \mathbf{y}$, das gleich einem zu den Vektoren $\mathbf{x}$ und $\mathbf{y}$ orthogonalen Vektor $\mathbf{z}$ ist, der damit einen *Normalenvektor* für die durch $\mathbf{x}$ und $\mathbf{y}$ aufgespannte Ebene bildet. Normalenvektoren erweisen sich u. a. als nützlich, wenn Teilräume von Vektorräumen charakterisiert werden sollen. Darüber hinaus liefert das vektorielle Produkt eine gewisse Veranschaulichung der Bedeutung von Determinanten von Matrizen, die in Abschnitt 3.7, Seite 83, besprochen werden. Das vektorielle Produkt spielt in der multivariaten Analyse keine bedeutende Rolle und wird hier mehr aus Vollständigkeitsgründen behandelt; es wird in Abschnitt 8.1 vorgestellt, der aber für das Verständnis der folgenden Abschnitte nicht wesentlich ist und bei der Lektüre übersprungen werden kann, wenn man sich nur über die für die hauptsächlichen Anwendungen der Vektor- und Matrixrechnung in der multivariaten Statistik interessiert.

Das Skalarprodukt Das Skalarprodukt wird in der folgenden Definition eingeführt:

Definition 2.5 Es seien $\mathbf{x} = (x_1, \dots, x_n)'$ und $\mathbf{y} = (y_1, \dots, y_n)'$ zwei n -dimensionale Vektoren. Dann ist das Skalarprodukt von $\mathbf{x}$ und $\mathbf{y}$ durch

$$\mathbf{x}'\mathbf{y} = \sum_{i=1}^n x_i y_i \quad (2.20)$$

definiert.

$\mathbf{x}'\mathbf{y}$ ist ein Skalar, also eine "einfache" Zahl (kein Vektor), deswegen der Ausdruck 'Skalarprodukt'. Die Schreibweise $\mathbf{x}'\mathbf{y}$ legt nahe, dass das Skalarprodukt in der Form

$$\mathbf{x}'\mathbf{y} = (x_1, x_2, \dots, x_n) \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix} = \sum_{i=1}^n x_i y_i \quad (2.21)$$

angeschrieben wird; der Sinn dieser Anordnung von $\mathbf{x}$ und $\mathbf{y}$ wird deutlich, wenn die Multiplikation von Matrizen erklärt wird. Statt des Ausdrucks 'Skalarprodukt' kommt auch der Ausdruck 'inneres Produkt' vor, und statt der Bezeichnung $\mathbf{x}'\mathbf{y}$ kommen auch die Schreibweisen $\mathbf{x}^T \mathbf{y}$, $\langle \mathbf{x}, \mathbf{y} \rangle$ und $\mathbf{x} \cdot \mathbf{y}$ vor; insbesondere in der anglo-amerikanischen Literatur findet man auch den Ausdruck *dot product* für das Skalarprodukt $\mathbf{x}'\mathbf{y}$ (der Ausdruck *scalar product* ist aber ebenfalls üblich).

Es zeigt sich (s. unten), dass der Mittelwert, die Varianz, die Kovarianz und der Produkt-Moment-Korrelationskoeffizient Spezialfälle von Skalarprodukten sind.

Länge und Normierung von Vektoren Für $\mathbf{x} = \mathbf{y}$ erhält man

$$\mathbf{x}'\mathbf{x} = \sum_{i=1}^n x_i^2 = \|\mathbf{x}\|^2 \quad (2.22)$$

Am Beispiel $n = 2$ sieht man, dass $\|\mathbf{x}\|^2$ das Quadrat der Länge des Vektors $\mathbf{x}$ ist (Satz des Pythagoras). Also ist $\|\mathbf{x}\| = \sqrt{\mathbf{x}'\mathbf{x}}$ die *Norm*, d.h. die Länge des Vektors.

Es sei $\vec{0} \neq \|\mathbf{x}\| \neq 1$. Dann kann $\mathbf{x}$ *normiert* werden. Dies geschieht durch Multiplikation mit einem geeignet gewählten Skalar λ :

$$\|\lambda\mathbf{x}\| = |\lambda|\|\mathbf{x}\| = 1.$$

Daraus folgt sofort

$$\lambda = \frac{1}{\|\mathbf{x}\|} > 0, \quad (2.23)$$

denn $\|\mathbf{x}\| > 0$. Ein Vektor mit der Länge $\|\mathbf{x}\|$ wird also normiert, indem man seine Komponenten mit $1/\|\mathbf{x}\|$ multipliziert.

Generell bedeutet die Multiplikation eines Vektors mit einem Skalar eine Skalierung der Länge des Vektors:

$$\|\lambda\mathbf{x}\|^2 = \sum_{i=1}^n (\lambda x_i)^2 = \lambda^2 \|\mathbf{x}\|^2,$$

also

$$\|\lambda\mathbf{x}\| = |\lambda|\|\mathbf{x}\|. \quad (2.24)$$

Dass hier $|\lambda|$ statt nur λ steht, folgt daraus, dass notwendig $\|\lambda\mathbf{x}\| \geq 0$ sein muß, unabhängig davon, ob λ größer oder kleiner als Null ist.

Eigenschaften des Skalarprodukts: Das Skalarprodukt hat die folgenden Eigenschaften bzw. für das Skalarprodukt gelten die folgenden Rechenregeln (alle Vektoren sind n -dimensional):

Positive Definitheit	$\mathbf{x}'\mathbf{x} = \langle \mathbf{x}, \mathbf{x} \rangle \geq 0, \mathbf{x}'\mathbf{x} = 0 \Leftrightarrow \mathbf{x} = \vec{0}.$	
Assoziativität	$\lambda \in \mathbb{R}, \langle \lambda\mathbf{x}, \mathbf{y} \rangle = \langle \mathbf{x}, \lambda\mathbf{y} \rangle = \lambda \langle \mathbf{x}, \mathbf{y} \rangle$	(2.25)
Kommutativität	$\mathbf{x}'\mathbf{y} = \mathbf{y}'\mathbf{x}$	
Distributivität	$(\mathbf{x} + \mathbf{y})'\mathbf{z} = \mathbf{x}'\mathbf{z} + \mathbf{y}'\mathbf{z}.$	

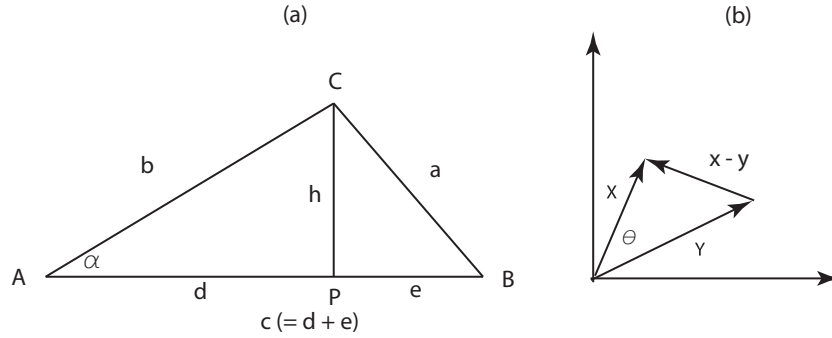
Diese Eigenschaften folgen sofort aus der Definition (2.20) des Skalarprodukts.

Das Skalarprodukt und der Winkel zwischen zwei Vektoren Es wird zunächst an den Kosinussatz erinnert: Es gilt (vergl. Abb. 1, (a))

$$a^2 = b^2 + c^2 - 2bc \cos \alpha. \quad (2.26)$$

Beweis: h ist das von Punkt C auf die Verbindungsline $c = \overline{AB}$ gefällte Lot

Abbildung 4: Zum Kosinussatz $a^2 = b^2 + c^2 - 2bc \cos \alpha$



(P). Es ist $d = \overline{AP}$, $e = \overline{PB}$. Nach dem Satz des Pythagoras ist $a^2 = h^2 + e^2$ und $b^2 = h^2 + d^2$, d.h. $h^2 = b^2 - d^2$, und nach Abb. 4 ist $e^2 = (c - d)^2$, so dass

$$a^2 = h^2 + e^2 = b^2 - d^2 + (c - d)^2 = b^2 + c^2 - 2cd$$

folgt. Weiter gilt $\cos \alpha = d/b$, dh $d = b \cos \alpha$. Damit erhält man $a^2 = b^2 + c^2 - 2bc \cos \alpha$. $\square$

Anmerkung: Für $\alpha = \pi/2$ (90°) ist (2.26) gerade die Aussage des Satzes von Pythagoras. Umgekehrt kann man (2.26) für $\alpha \neq \pi/2$ als Verallgemeinerung des Satzes des Pythagoras ansehen. $\square$

In vektorieller Schreibweise nimmt der Kosinussatz (2.26) die Form

$$\|\mathbf{x} - \mathbf{y}\|^2 = \|\mathbf{x}\|^2 + \|\mathbf{y}\|^2 - 2\|\mathbf{x}\|\|\mathbf{y}\|\cos \alpha \quad (2.27)$$

an (vergl. Abbildung 4 (b)). Für $\alpha = \pi/2$, also für einen Winkel von 90° , folgt $\cos \alpha = 0$ und es ergibt sich der Satz des Pythagoras in Vektorschreibweise. Hieraus folgt eine Beziehung zwischen dem Skalarprodukt $\mathbf{x}'\mathbf{y}$ und dem Kosinus des Winkels α : es ist

$$\|\mathbf{x} - \mathbf{y}\|^2 = \sum_i (x_i - y_i)^2 = \sum_i x_i^2 + \sum_i y_i^2 - 2 \sum_i x_i y_i = \|\mathbf{x}\|^2 + \|\mathbf{y}\|^2 - 2\mathbf{x}'\mathbf{y}.$$

Setzt man diesen Ausdruck für $\|\mathbf{x} - \mathbf{y}\|^2$ in (2.27) ein, so wird man auf die Beziehung

$$\mathbf{x}'\mathbf{y} = \|\mathbf{x}\|\|\mathbf{y}\|\cos \alpha \quad (2.28)$$

geführt. Das Skalarprodukt wird also einerseits durch die Längen $\|\mathbf{x}\|, \|\mathbf{y}\|$ der Vektoren und andererseits durch den Winkel α zwischen den Vektoren bestimmt. Für einen gegebenen Winkel α kann man den Effekt der Länge leicht ausdrücken: Die Länge eines Vektors, etwa des Vektors $\mathbf{x}$, werde um den Faktor λ verändert, so dass aus $\mathbf{x}$ der Vektor $\tilde{\mathbf{x}} = \lambda\mathbf{x}$ entsteht. Dann ist

$$\tilde{\mathbf{x}}'\mathbf{y} = \lambda\|\mathbf{x}\|\|\mathbf{y}\|\cos \alpha = \lambda\mathbf{x}'\mathbf{y}, \quad (2.29)$$

d.h. des Skalarprodukt wird dann um den gleichen Faktor λ vergrößert ($\lambda > 1$) oder verkleinert ($\lambda < 1$).

Man macht sich leicht klar, in welcher Weise das Skalarprodukt für gegebene Vektorlängen $\|\mathbf{x}\|$ und $\|\mathbf{y}\|$ vom Winkel α zwischen den Vektoren abhängt. Man hat

$$\mathbf{x}'\mathbf{y} = \begin{cases} \|\mathbf{x}\|\|\mathbf{y}\|, & \alpha = 0, \quad \cos \alpha = 1 \\ 0, & \alpha = \frac{\pi}{2}, \quad \cos \alpha = 0 \\ -\|\mathbf{x}\|\|\mathbf{y}\|, & \alpha = \pi, \quad \cos \alpha = -1 \end{cases} \quad (2.30)$$

Für $\alpha = 0$ wird $\mathbf{x}'\mathbf{y}$ also maximal, für $\alpha = \pi$ wird $\mathbf{x}'\mathbf{y}$ minimal, und für $\alpha = \pi/2$ nimmt $\mathbf{x}'\mathbf{y}$ den Wert Null an.

Man kann die Gleichung (2.28) nach $\cos \alpha$ auflösen:

$$\cos \alpha = \frac{\mathbf{x}'\mathbf{y}}{\|\mathbf{x}\|\|\mathbf{y}\|} = \frac{\mathbf{x}'}{\|\mathbf{x}\|} \frac{\mathbf{y}}{\|\mathbf{y}\|} \quad (2.31)$$

Wie die rechte Seite zeigt, ist $\cos \alpha$ gleich dem Skalarprodukt der normierten Vektoren $\mathbf{x}/\|\mathbf{x}\|$ und $\mathbf{y}/\|\mathbf{y}\|$; dies bedeutet eine Normierung des Skalarprodukts auf das Intervall $[-1, +1]$, denn anhand der vorausgegangenen Gleichungen verifiziert man leicht

$$\frac{\mathbf{x}'\mathbf{y}}{\|\mathbf{x}\|\|\mathbf{y}\|} = \begin{cases} 1, & \alpha = 0 \\ 0, & \alpha = \frac{\pi}{2}, \\ -1, & \alpha = \pi. \end{cases} \quad (2.32)$$

Für $\alpha = \pi/2$ ($= 90^\circ$) ist $\cos \alpha = 0$, dann folgt $\mathbf{x}'\mathbf{y} = 0$. Für den Fall $\alpha = \pi/2$ stehen die Vektoren $\mathbf{x}$ und $\mathbf{y}$ senkrecht aufeinander, d.h. sie bilden einen rechten Winkel. Deshalb heißen die beiden Vektoren *orthogonal*⁶ zueinander.

Anmerkungen:

1. Für zwei dreidimensionale Vektoren ist den Winkel zwischen ihnen ebenfalls definiert. Im allgemeinen n -dimensionalen Fall mit $n > 3$ ist intuitiv nicht mehr klar, was ein Winkel zwischen den Vektoren bedeuten soll; in diesem Fall wird die rechte Seite von (2.31) als *Definition* des Winkels zwischen ihnen interpretiert.
2. Korrespondierend zur Definition von $\cos \theta$ durch (2.31) wird gelegentlich (2.28) als *Definition* des Skalarprodukts $\mathbf{x}'\mathbf{y}$ verwendet, – aus der dann der Kosinussatz mit dem Spezialfall der Satz des Pythagoras und letztlich die ursprüngliche Definition (2.20) des Skalarprodukts *folgt*.

⁶von griechisch orthos (ὀρθός) = richtig, recht-, und gonia (γωνία) = Ecke, Winkel

Beispiel 2.2 Der Mittelwert als Skalarprodukt Es sei $\mathbf{x} = (x_1, x_2, \dots, x_n)'$ ein Vektor mit n Messwerten, und $\vec{1}$ sei der in (2.11) eingeführte n -dimensionale Einsvektor. Dann ist die Summe der Komponenten durch $\mathbf{x}'\vec{1}$ gegeben und ihr Mittelwert läßt sich durch

$$\bar{x} = \frac{1}{n} \vec{1}' \mathbf{x} = \frac{1}{n} \mathbf{x}' \vec{1} \quad (2.33)$$

ausdrücken. Im Übrigen ist $\vec{1}'\vec{1} = \|\vec{1}\|^2 = n$, so dass die Länge von $\vec{1}$ durch $\|\vec{1}\| = \sqrt{n}$ gegeben ist. $\square$

Beispiel 2.3 Varianz und Standardabweichung Es sei $\mathbf{X}$ ein n -dimensionaler Vektor von Messwerten. Der Mittelwert der Messwerte, also der Komponenten von $\mathbf{X}$, ist nach (2.33) durch $\bar{x} = \frac{1}{n} \vec{1}' \mathbf{x}$ gegeben. Dann ist

$$\mathbf{x} = \mathbf{X} - \frac{1}{n} \vec{1}' \mathbf{x} \quad (2.34)$$

der Vektor der Abweichungen der Komponenten von $\mathbf{X}$ vom Mittelwert ($\vec{1}' \mathbf{x} = \bar{x}$ ist ein Skalar, mit dem der Vektor $\vec{1}$ multipliziert wird, $\vec{1}' \mathbf{x} \vec{1}$ ist also ein Vektor, dessen Komponenten alle gleich $\bar{x}$ sind). Dann ist $\|\mathbf{x}\|^2$ die Summe der Quadrate der Abweichungen vom Mittelwert, so dass

$$s^2 = \frac{1}{n} \|\mathbf{x}\|^2 \quad (2.35)$$

ein Ausdruck für die Stichprobenvarianz der Messwerte ist (natürlich kann man auch durch $n-1$ teilen, um den Bias dieser Varianzschätzung auszugleichen), und

$$s = \frac{1}{\sqrt{n}} \|\mathbf{x}\| \quad (2.36)$$

ist ein Ausdruck für die Standardabweichung. Unter der Voraussetzung, dass $\mathbf{x}$ wie in (2.34) definiert ist, ist $\|\mathbf{x}\|^2$ ist also proportional zu Varianz, und die Länge des Vektors $\mathbf{x}$ ist proportional zur Standardabweichung der Messwerte. $\square$

Beispiel 2.4 Der Korrelationskoeffizient Es seien X_i und Y_i Messwerte mit den Mittelwerten $\bar{x}$ bzw. $\bar{y}$, und es sei $x_i = X_i - \bar{x}$ bzw. $y_i = Y_i - \bar{y}$, $i = 1, \dots, n$. Dann ist⁷

$$\frac{1}{n} \mathbf{x}' \mathbf{y} = \frac{1}{n} \sum_{i=1}^n x_i y_i$$

die Kovarianz der Messwerte und

$$\frac{1}{n} \|\mathbf{x}\|^2 = \frac{1}{n} \sum_{i=1}^n x_i^2 = s_x^2, \quad \frac{1}{n} \|\mathbf{y}\|^2 = \frac{1}{n} \sum_{i=1}^n y_i^2 = s_y^2$$

⁷Üblicherweise wird durch $n-1$ statt durch n geteilt, um eine Verzerrung der Schätzungen von Varianz und Kovarianz auszugleichen. Dieser Effekt kann hier vernachlässigt werden, da sich der Faktor $1/n$ bzw. $1/(n-1)$ herauskürzt.

sind die Varianzen der Messwerte. Dann ist

$$r_{xy} = \frac{\mathbf{x}'\mathbf{y}}{\|\mathbf{x}\|\|\mathbf{y}\|} = \cos \theta \quad (2.37)$$

gleich dem Korrelationskoeffizienten für X und Y , d.h. $r_{xy} = \cos \theta$, θ der Winkel zwischen $\mathbf{x}$ und $\mathbf{y}$; auf diesen Zusammenhang wird bei der Faktorenanalyse zurückgegriffen. Der Maximalwert von $\cos \theta$ ist $+1$, der Minimalwert ist -1 , so dass

$$-1 \leq r_{xy} \leq 1. \quad (2.38)$$

□

Skalarprodukt und Ähnlichkeit: Das Skalarprodukt kann zur Definition eines Ähnlichkeitsmaßes für die Objekte x und y , die durch die Vektoren $\mathbf{x}$ und $\mathbf{y}$ repräsentiert werden verwendet werden. Ähnlichkeit wird üblicherweise durch eine Maßzahl $\mathbf{s}$ zwischen Null und Eins abgebildet, $0 \leq \mathbf{s}(x, y) \leq 1$. $\mathbf{s} = 1$ steht für $\mathbf{x} = \mathbf{y}$ (die Objekte x und y müssen deswegen nicht identisch sein!), und $\mathbf{s} = 0$ steht für vollkommene Unähnlichkeit, und im Prinzip kann $\mathbf{s}(x, y) \neq \mathbf{s}(y, x)$ sein. $\mathbf{s}$ kann auf sehr verschiedene Weise definiert werden; das Skalarprodukt ist eine spezielle Definition: Es wird

$$\mathbf{s}(\mathbf{x}, \mathbf{y}) = \cos \alpha(\mathbf{x}, \mathbf{y}) = \frac{\mathbf{x}'\mathbf{y}}{\|\mathbf{x}\|\|\mathbf{y}\|} \quad (2.39)$$

gesetzt. Man könnte auch einfach $\mathbf{s} = \mathbf{x}'\mathbf{y}$ setzen, hätte dann aber kein normiertes (Ähnlichkeits-)Maß für $\mathbf{s}$. Hier wird $\mathbf{s}$ offenbar als symmetrisches Maß konzipiert. Für $\alpha = 0$ wird die Ähnlichkeit maximal, und für $\alpha = \pi/2$, wenn $\mathbf{x}$ und $\mathbf{y}$ also orthogonal zueinander sind, wird die Ähnlichkeit minimal, nämlich gleich Null. Vollständige Unähnlichkeit bedeutet also nicht notwendig, dass x und y keine gemeinsamen Merkmale haben, sondern nur, dass sie durch orthogonale Vektoren repräsentiert werden.

Skalarprodukte von Messwerten und zentrierten Messwerten: Gegeben seien zwei Reihen von Messwerten: $X_1, X_2, \dots, X_n$ und $Y_1, Y_2, \dots, Y_n$. Man kann sie als Vektoren $\mathbf{X}$ und $\mathbf{Y}$ auffassen und, um die "Ähnlichkeit" der Messwerte zu bestimmen, das Skalarprodukt $\mathbf{x}'\mathbf{y}$ berechnen. Ebenso kann man die Kovarianz als Ähnlichkeitsmaß auffassen, oder auch den Korrelationskoeffizienten. Die Berechnung der Kovarianz setzt voraus, dass die Messwerte *zentriert* werden, d.h. dass man von den X_i zu den $x_i = X_i - \bar{x}$ und von den Y_i zu den $y_i = Y_i - \bar{y}$ übergeht. Die Skalarprodukte von unzentrierten und zentrierten Werten unterscheiden sich natürlich:

$$\mathbf{x}'\mathbf{y} = \sum_{i=1}^n x_i y_i = \sum_{i=1}^n (X_i - \bar{x})(Y_i - \bar{y}) = \sum_{i=1}^n X_i Y_i - n\bar{x}\bar{y} = \mathbf{X}'\mathbf{Y} - n\bar{x}\bar{y}. \quad (2.40)$$

Auch der Winkel zwischen unzentrierten und zentrierten Vektoren wird im Allgemeinen verschieden sein: (2.40) impliziert

$$\mathbf{x}'\mathbf{x} = \|\mathbf{x}\|^2 = \|\mathbf{X}\|^2 - n\bar{x}^2, \quad \mathbf{y}'\mathbf{y} = \|\mathbf{Y}\|^2 - n\bar{y}^2, \quad (2.41)$$

so dass

$$\mathbf{x}'\mathbf{y} = \frac{\mathbf{X}'\mathbf{Y} - n\bar{x}\bar{y}}{\sqrt{\|\mathbf{X}\|^2 - n\bar{x}^2} \sqrt{\|\mathbf{Y}\|^2 - n\bar{y}^2}} = \cos \alpha_{xy}, \quad (2.42)$$

während

$$\mathbf{X}'\mathbf{Y} = \frac{\mathbf{X}'\mathbf{Y}}{\|\mathbf{X}\| \|\mathbf{Y}\|} = \cos \alpha, \quad (2.43)$$

und der Vergleich von (2.42) und (2.43) zeigt, dass der Fall $\alpha = \alpha_{xy}$ allenfalls in speziellen Fällen vorliegt.

Beispiel 2.5 Einheitsvektoren sind orthogonal Die Einheitsvektoren $\mathbf{e}_j$ und $\mathbf{e}_k$, $j \neq k$, sind orthogonal. Denn

$$\mathbf{e}_j'\mathbf{e}_k = 0 \cdot 0 + \dots + 1 \cdot 0 + \dots + 0 \cdot 1 + 0 \cdot 0 + \dots + 0 \cdot 0 = 0, \quad (2.44)$$

wobei in $1 \cdot 0$ die 1 der j -ten Komponente von $\mathbf{e}_j$ und in $0 \cdot 1$ die 1 an der k -ten Stelle von $\mathbf{e}_k$ gemeint ist. Im Übrigen sieht man leicht, dass $\|\mathbf{e}_j\| = 1$ für alle $\mathbf{e}_j$. $\square$

Beispiel 2.6 Multiple Regression Für zentrierte Werte nimmt die multiple Regression die Form

$$\mathbf{y} = b_1\mathbf{x}_1 + \dots + b_p\mathbf{x}_p + \mathbf{e} \quad (2.45)$$

an. Setzt man

$$\hat{\mathbf{y}} = b_1\mathbf{x}_1 + \dots + b_p\mathbf{x}_p,$$

so hat man für die i -te Komponente von $\hat{\mathbf{y}}$ den Ausdruck

$$y_i = b_1x_{i1} + \dots + b_px_{ip}, \quad (2.46)$$

d.h. $\hat{y}_i$ ist das Skalarprodukt der Vektoren $\mathbf{x}_i = (x_{i1}, \dots, x_{ip})'$ und $\mathbf{b} = (b_1, \dots, b_p)'$, so dass man

$$\hat{y}_i = \mathbf{x}_i'\mathbf{b} = \|\mathbf{x}_i\| \|\mathbf{b}\| \cos \alpha_{ib}. \quad (2.47)$$

wobei α_{ib} der Winkel zwischen $\mathbf{x}_i$ und $\mathbf{b}$ ist. Für eine gegebene Länge $\|\mathbf{b}\|$ des Parametervektors $\mathbf{b}$ und einen gegebenen Winkel $\alpha_{ib} \neq \pi/2$ wird das Skalarprodukt $\hat{y}_i$ um so größer, je länger $\mathbf{x}_i$ ist. Für gegebene Längen $\|\mathbf{x}_i\|$ und $\|\mathbf{b}\|$ wird $\hat{y}_i$ um so größer, je kleiner der Winkel α_{ib} ist, und es wird maximal, wenn $\alpha_{ib} = 0$, wenn also $\mathbf{x}_i$ und $\mathbf{b}$ dieselbe Orientierung haben. Die Vektoren $\mathbf{x}_i$ und $\mathbf{b}$ unterscheiden sich dann nur durch ihre Längen, d.h. $\mathbf{x}_i = \lambda\mathbf{b}$, die Komponenten x_{ij} des Vektors $\mathbf{x}_i$ unterscheiden sich dann nur durch einen Proportionalitätsfaktor λ von den korrespondierenden Komponenten b_j des "Gewichts"vektors $\mathbf{b}$.

$\hat{y}_i = 0$, wenn $\hat{Y}_i = \bar{\hat{y}}$, wenn der vorhergesagte Rohwert Y_i also gerade gleich dem durchschnittlichen vorhergesagten Rohwert ist. Das ist einerseits der Fall, wenn $\mathbf{x}_i = \vec{0}$, d.h. wenn die Prädiktorwerte x_i gerade gleich den entsprechenden Mittelwerten sind. Andererseits ist für $\mathbf{x}_i \neq \vec{0}$ der vorhergesagte Wert $\hat{y}_i = 0$, wenn das Skalarprodukt $\mathbf{x}_i \mathbf{b} = 0$ ist, wenn also $\mathbf{b}$ und $\mathbf{x}_i$ orthogonal sind. Die Prädiktorwerte (die Komponenten von $\mathbf{x}_i$, z.B. die Ausprägungen von Symptomen) korrespondieren dann nicht zu der für eine korrekte Vorhersage der Kriteriumsvariablen Y notwendigen Art der Zusammensetzung (Gewichtung). $\square$

Beispiel 2.7 Testscores und Skalarprodukt Nach dem faktorenanalytischen Testmodell ergibt sich z.B. der Testscore x_i eines Probanden P_i gemäß

$$x_i = a_1 F_{i1} + a_2 F_{i2} + \cdots + a_r F_{ir} + e_i = \hat{x}_i + e_i, \quad (2.48)$$

$$\hat{x}_i = a_1 F_{i1} + a_2 F_{i2} + \cdots + a_r F_{ir} \quad (2.49)$$

wobei e_i den unvermeidlichen Fehler repräsentiert. Die F_{ik} , $k = 1, \dots, r$ repräsentieren das jeweilige Ausmaß des Merkmals auf der k -ten latenten Dimension. Die Koeffizienten $a_1, \dots, a_r$ erweisen sich unter der Bedingung, dass die latenten Variablen unabhängig voneinander sind, als Korrelationen zwischen dem Test und den latenten Merkmalen. Faßt man die a_k und die F_{ik} als Komponenten von Vektoren $\mathbf{a} = (a_1, \dots, a_r)'$ und $\mathbf{F}^i = (F_{i1}, \dots, F_{ir})'$ auf, so besagt das Modell⁸ (2.48), dass der Messwert x_i gerade als Skalarprodukt von $\mathbf{a}$ und $\mathbf{F}^i$ definiert ist:

$$x_i = \mathbf{a}' \mathbf{F}^i = \|\mathbf{a}\| \|\mathbf{F}^i\| \cos \theta, \quad (2.50)$$

wobei von (2.28) Gebrauch gemacht wurde; θ ist der Winkel zwischen dem Vektor $\mathbf{a}$ und dem Vektor $\mathbf{F}^i$. Offenbar ist $x_i = 0$ dann, wenn $\theta = \pi/2$ ist, wenn also die Vektoren $\mathbf{a}$ und $\mathbf{F}^i$ orthogonal sind, wenn also die Ausprägungen der oder des Probanden auf den latenten Dimensionen gewissermaßen nicht mit der Gewichtung, mit der die latenten Merkmale in den Test eingehen, korrelieren.

Man kann man die Länge $\|\mathbf{a}\|$ als Gesamtmaß für das Erfassen der latenten Dimensionen durch den Test und die Länge $\|\mathbf{F}^i\|$ als Gesamtmaß für die Ausprägung des gemessenen Merkmals von P_i definieren. $\square$

Es gelte insbesondere $\mathbf{y} = \lambda \mathbf{x}$, λ ein Skalar, d.h. die beiden Vektoren haben die gleiche Orientierung, aber für $\lambda \neq 1$ verschiedene Längen. Aus (2.31) folgt dann

$$\cos \theta = \frac{\lambda \|\mathbf{x}\|^2}{\lambda \|\mathbf{x}\|^2} = 1.$$

⁸Die Gleichung (2.48) repräsentiert ein *Modell* insofern, als sie eine mögliche Annahme über das Zustandekommen eines Messwertes darstellt. Andere Annahmen sind ebenfalls denkbar, z.B. $x_i = a_1 F_{i1} + a_2 F_{i2} + a_3 F_{i1} F_{i2}$; hier werden nur zwei latente Dimensionen angenommen, deren Wechselwirkung in Form des Produkts $F_{i1} F_{i2}$ ebenfalls den Wert von x_i bestimmen. Solche Modelle werden in diesem Skriptum nicht betrachtet.

was natürlich $\theta = 0$ bedeutet, und aus (2.53) folgt dann, dass für $\mathbf{y} = \lambda \mathbf{x}$, also für

$$\lambda x_i = y_i, \quad i = 1, \dots, n \quad (2.51)$$

das Skalarprodukt $\mathbf{x}'\mathbf{y}$ den maximal möglichen Wert annimmt, d.h. es gilt dann

$$\mathbf{x}'\mathbf{y} = \|\mathbf{x}\|\|\mathbf{y}\|, \quad \text{wenn } \mathbf{y} = \lambda \mathbf{x} \quad (2.52)$$

2.2.4 Die Cauchy-Schwarzsche Ungleichung

Da $\cos \theta \leq 1$ folgt aus

$$\begin{aligned} \mathbf{x}'\mathbf{y} &= \|\mathbf{x}\|\|\mathbf{y}\| \cos \theta \\ \mathbf{x}'\mathbf{y} &\leq \|\mathbf{x}\|\|\mathbf{y}\|. \end{aligned} \quad (2.53)$$

Quadriert man beide Seiten der Ungleichung, so erhält man

$$|\mathbf{x}'\mathbf{y}|^2 \leq \|\mathbf{x}\|^2 \|\mathbf{y}\|^2; \quad (2.54)$$

oder

$$|\mathbf{x}'\mathbf{y}| \leq \|\mathbf{x}\|\|\mathbf{y}\|; \quad (2.55)$$

Diese Ungleichung ist die *Cauchy-Schwarzsche Ungleichung*. Denn aus $\|\mathbf{x}\|\|\mathbf{y}\| \cos \theta = \mathbf{x}'\mathbf{y}$ folgt

$$\|\mathbf{x}\|^2 \|\mathbf{y}\|^2 \cos^2 \theta = (\mathbf{x}'\mathbf{y})^2 \quad (2.56)$$

und $\cos^2 \theta \leq 1$ impliziert (2.54) bzw. (2.55). $\theta = 0$ bedeutet, dass $\mathbf{x}$ und $\mathbf{y}$ parallel sind. In diesem Fall ist $\cos \theta = 1$ und (2.56) impliziert

$$\|\mathbf{x}\|^2 \|\mathbf{y}\|^2 = (\mathbf{x}'\mathbf{y})^2 \Rightarrow \|\mathbf{x}\|\|\mathbf{y}\| = |\mathbf{x}'\mathbf{y}|, \quad (2.57)$$

und umgekehrt impliziert $\theta = 0$ die Beziehung (2.57).

Beispiel 2.8 Der Korrelationskoeffizient In Gleichung (2.37) wurde der Korrelationskoeffizient in vektorieller Schreibweise gegeben:

$$r_{xy} = \frac{\mathbf{x}'\mathbf{y}}{\|\mathbf{x}\|\|\mathbf{y}\|}.$$

Wegen (2.56) folgt sofort $|r_{xy}| \leq 1$. □

Beispiel 2.9 Fortsetzung von Beispiel 2.7 In Gleichung (2.50), Seite 27, wurde ein Messwert x_i im Rahmen eines faktorenanalytischen Modells als Skalarprodukt von einem Vektor $\mathbf{a}$ von "Gewichten" und einem Vektor $\mathbf{F}^i$ von Ausprägungen auf latenten Dimensionen betrachtet. Nach (2.55) und (2.52) nimmt dann x_i den maximal möglichen Wert $\|\mathbf{a}\|\|\mathbf{F}^i\|$ an, wenn $\mathbf{F}^i = \lambda \mathbf{a}$ gilt, wenn also das Profil von P_i dem Profil $\mathbf{a}$ des Tests bis auf eine Proportionalitätskonstante gleicht. Sind die Vektoren $\mathbf{a}$ und $\mathbf{F}^i$ aber orthogonal, so ist der Testwert $x_i = 0$; die Anteile $a_1, \dots, a_r$, mit denen der Test die latenten Dimensionen erfaßt, korrelieren gewissermaßen nicht mit dem Muster der Ausprägungen der latenten Dimensionen bei der i -ten Person. Dies gilt unabhängig vom Wert $\|\mathbf{F}^i\|$, der Länge des Vektors $\mathbf{F}^i$, der die Gesamtbegabung des i -ten Probanden abbildet. □

2.2.5 Das dyadische Produkt $\mathbf{xy}'$

Das in der folgenden Definition charakterisierte Produkt zweier Vektoren ergibt weder einen Skalar noch einen Vektor, sondern eine Matrix:

Definition 2.6 Es seien $\mathbf{x}$ und $\mathbf{y}$ zwei Vektoren. Das durch die Matrix

$$\mathbf{xy}' = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} (y_1, y_2, \dots, y_m) = \begin{pmatrix} x_1 y_1 & x_1 y_2 & \cdots & x_1 y_m \\ x_2 y_1 & x_2 y_2 & \cdots & x_2 y_m \\ & & \ddots & \\ x_n y_1 & x_n y_2 & \cdots & x_n y_m \end{pmatrix} \quad (2.58)$$

definierte Produkt $\mathbf{x}'\mathbf{y}$ heißt dyadisches Produkt der Vektoren $\mathbf{x}$ und $\mathbf{y}$.

Anmerkung: Gelegentlich wird das dyadische Produkt auch *äußeres Produkt* genannt, um es vom Skalarprodukt als *inneren Produkt* zu unterscheiden. Das kann gelegentlich verwirrend sein, da der Ausdruck 'äußeres Produkt' ja auch für das vektorielle Produkt verwendet wird. $\square$

Das dyadische Produkt findet im Rahmen der multivariaten Statistik hauptsächlich Anwendung bei der Definition und Interpretation von Varianz-Kovarianz-Matrizen, sowie bei der Spektraldarstellung symmetrischer Matrizen (Abschnitt 3.9.4, Gleichung (3.149), Seite 111).

2.3 Lineare Unabhängigkeit von Vektoren

2.3.1 Definition der linearen Unabhängigkeit

Viele Probleme der multivariaten Statistik führen auf die Frage, ob gegebene Vektoren als Linearkombinationen anderer Vektoren darstellbar sind, und wenn ja, wieviele dieser anderen Vektoren dazu maximal notwendig sind. Andere Probleme führen auf die Frage, ob bestimmte Gleichungssysteme lösbar sind, und wenn ja, ob die Lösung eindeutig ist oder nicht. Diese Fragen lassen sich zusammenfassend diskutieren, wenn man den Begriff der linearen Abhängigkeit bzw. Unabhängigkeit von Vektoren berücksichtigt.

Zur Einführung und Motivation werde der übersichtliche Fall von beliebigen 2-dimensionalen Vektoren $\mathbf{y}$, $\mathbf{x}_1$ und $\mathbf{x}_2$ betrachtet. Die Frage sei, ob sich $\mathbf{y}$ als Linearkombination der Vektoren $\mathbf{x}_1$ und $\mathbf{x}_2$ darstellen läßt, d.h. ob Koeffizienten $a_1, a_2 \in \mathbb{R}$ existieren derart, dass

$$\mathbf{y} = a_1 \mathbf{x}_1 + a_2 \mathbf{x}_2 = a_1 \begin{pmatrix} x_{11} \\ x_{21} \end{pmatrix} + a_2 \begin{pmatrix} x_{12} \\ x_{22} \end{pmatrix} \quad (2.59)$$

gilt. Schreibt man die rechte Seite aus ergibt sich das System von Gleichungen

$$y_1 = a_1 x_{11} + a_2 x_{12} \quad (2.60)$$

$$y_2 = a_1 x_{21} + a_2 x_{22} \quad (2.61)$$

mit $\mathbf{y} = (y_1, y_2)'$, $\mathbf{x}_1 = (x_{11}, x_{21})'$, $\mathbf{x}_2 = (x_{12}, x_{22})'$. Man findet

$$a_1 = \frac{y_1 x_{22} - y_2 x_{12}}{x_{11} x_{22} - x_{12} x_{21}} \quad (2.62)$$

$$a_2 = \frac{y_2 x_{11} - y_1 x_{21}}{x_{11} x_{22} - x_{12} x_{21}} \quad (2.63)$$

Es wird deutlich, dass eine *notwendige Bedingung* für die Existenz einer Lösung $\mathbf{a} = (a_1, a_2)'$ durch

$$x_{11} x_{22} - x_{12} x_{21} \neq 0 \quad (2.64)$$

gegeben ist. Denn $x_{11} x_{22} - x_{12} x_{21} = 0$ würde bedeuten, dass durch 0 dividiert werden muß, damit man eine Lösung erhält, – und diese Operation macht bekanntlich keinen Sinn.

Nun betrachte man den Fall

$$x_{11} x_{22} - x_{12} x_{21} = 0. \quad (2.65)$$

Er impliziert

$$\frac{x_{21}}{x_{11}} = \frac{x_{22}}{x_{12}} \quad (2.66)$$

Aber $x_{21}/x_{11} = \tan \theta$, θ der Winkel, der die Orientierung von $\mathbf{x}_1$ angibt (vergl. (2.7), Seite 13), und (2.66) besagt, dass dieser Winkel auch die Orientierung von $\mathbf{x}_2$ definiert. Daraus folgt, dass (2.65) dann erfüllt ist, wenn $\mathbf{x}_1$ und $\mathbf{x}_2$ parallel sind, und das heißt eben, dass sie dieselbe Orientierung haben. Dann gilt $\mathbf{x}_2 = \lambda \mathbf{x}_1$ und die Gleichung (2.59) geht über in die Gleichung

$$\mathbf{y} = a_1 \mathbf{x}_1 + a_2 \mathbf{x}_2 = a_1 \mathbf{x}_1 + a_2 \lambda \mathbf{x}_1 = (a_1 + a_2 \lambda) \mathbf{x}_1,$$

aus der sofort hervorgeht, dass (2.59) nur dann eine Lösung hat, wenn $\mathbf{y}$ tatsächlich dieselbe Orientierung wie $\mathbf{x}_1$ hat. Diese kleine Betrachtung illustriert noch einmal die Tatsache, dass es nicht möglich ist, aus zwei Vektoren mit identischer Orientierung einen Vektor zu erzeugen, der eine andere Orientierung hat.

Nun seien in (2.59) $\mathbf{x}_1$ und $\mathbf{x}_2$ nicht parallel. Die Frage ist, ob die Lösung (a_1, a_2) eindeutig ist. Dazu nehme man an, dass es eine zweite Lösung (b_1, b_2) gibt, so dass $\mathbf{y} = b_1 \mathbf{x}_1 + b_2 \mathbf{x}_2$ gilt. Dann folgt

$$\vec{0} = (a_1 - b_1) \mathbf{x}_1 + (a_2 - b_2) \mathbf{x}_2 = c_1 \mathbf{x}_1 + c_2 \mathbf{x}_2,$$

also $c_i = a_i - b_i$, $i = 1, 2$, so dass

$$\mathbf{x}_2 = \frac{c_2}{c_1} \mathbf{x}_1,$$

d.h. für $c_i \neq 0$, $i = 1, 2$, würde folgen, dass $\mathbf{x}_1$ und $\mathbf{x}_2$ entgegen der Voraussetzung dieselbe Orientierung haben. Für nicht parallele $\mathbf{x}_1, \mathbf{x}_2$ folgt also $c_i = a_i - b_i = 0$, also $a_i = b_i$ für alle i , d.h. die Lösung ist eindeutig.

Man kann die Diskussion zusammenfassen. Subtrahiert man auf beiden Seiten von (2.59) den Vektor $\mathbf{y}$, so erhält man

$$\vec{0} = a_1\mathbf{x}_1 + a_2\mathbf{x}_2 - \mathbf{y} = a_1\mathbf{x}_1 + a_2\mathbf{x}_2 + a_3\mathbf{y}, \quad a_3 = -1. \quad (2.67)$$

Ist der beliebig gewählte Vektor $\mathbf{y}$ als Linearkombination von $\mathbf{x}_1$ und $\mathbf{x}_2$ darstellbar, so sind offenbar nicht alle a_i in (2.67) gleich Null. Ist dagegen $\mathbf{y}$ nicht als Linearkombination von $\mathbf{x}_1$ und $\mathbf{x}_2$ darstellbar, so folgt, dass "nicht alle a_i sind gleich Null" eben nicht gilt⁹, d.h. aber, dass dann $a_i = 0$ für alle i gelten muß. Im ersten Fall kann man sagen, dass $\mathbf{y}$ von $\mathbf{x}_1$ und $\mathbf{x}_2$ *linear abhängig* ist, im zweiten Fall ist $\mathbf{y}$ nicht von $\mathbf{x}_1$ und $\mathbf{x}_2$ linear abhängig, d.h. $\mathbf{y}$ ist von diesen Vektoren *linear unabhängig*.

Diese Betrachtungen können auf den Fall m -dimensionaler Vektoren, m beliebig, übertragen werden. Gegeben seien n m -dimensionale Vektoren $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$, mit $\mathbf{x}_j \neq \vec{0}$ für alle j . Mindestens einer von ihnen sei als Linearkombination der übrigen Vektoren darstellbar; dann heißen die $\mathbf{x}_j$ *linear abhängig*. Ist keiner von ihnen als Linearkombination der restlichen Vektoren darstellbar, so heißen die $\mathbf{x}_j$ *linear unabhängig*.

Die Vektoren seien linear abhängig; da die Nummerierung der Vektoren beliebig vorgenommen werden kann, sei dann etwa $\mathbf{x}_1$ als Linearkombination der übrigen darstellbar, d.h. es existieren reelle Zahlen $\lambda_2, \dots, \lambda_n$, die *nicht alle* gleich Null sind, derart, dass

$$\mathbf{x}_1 = \lambda_2\mathbf{x}_2 + \dots + \lambda_n\mathbf{x}_n,$$

Diese Gleichung kann in der Form

$$\vec{0} = \lambda_1\mathbf{x}_1 + \lambda_2\mathbf{x}_2 + \dots + \lambda_n\mathbf{x}_n, \quad \lambda_1 = -1$$

geschrieben werden. Diese Gleichung ist eine Darstellung des Nullvektors als Linearkombination der $\mathbf{x}_j$. Nun werde angenommen, dass die $\mathbf{x}_j$ *nicht* linear abhängig sind. Dann folgt, dass die Aussage, nicht alle λ_j seien gleich Null, *nicht* gilt, d.h. es muß $\lambda_j = 0$ für alle $j = 1, \dots, n$ gelten. Für viele Betrachtungen ist es nützlich, von dieser formalen Charakterisierung der linearen Unabhängigkeit Gebrauch machen zu können, weshalb der Begriff der linearen Unabhängigkeit noch einmal explizit definiert wird:

Definition 2.7 Gegeben seien n m -dimensionale Vektoren $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$, mit $\mathbf{x}_j \neq \vec{0}$ für alle j . Die Vektoren heißen *linear unabhängig*, wenn

$$\vec{0} = \lambda_1\mathbf{x}_1 + \lambda_2\mathbf{x}_2 + \dots + \lambda_n\mathbf{x}_n \quad (2.68)$$

dann und nur dann gilt, wenn $\lambda_j = 0$ für alle $j = 1, \dots, n$.

⁹Dieser Schluß ist einfach eine Anwendung des *modus tollens*, demzufolge aus der Aussage $P \rightarrow Q$ folgt, dass $\neg Q \rightarrow \neg P$ folgt.

$\lambda_1 = \dots = \lambda_n = 0$ ist stets eine Lösung für (2.68), aber im Fall der linearen Unabhängigkeit der $\mathbf{x}_j$ ist dies die *einzigste* Lösung, und im Fall der linearen Abhängigkeit der $\mathbf{x}_j$ ist es *nicht die einzige* Lösung.

Satz 2.1 *Es seien $\mathbf{x}_1, \dots, \mathbf{x}_n$ linear unabhängige, m -dimensionale Vektoren, und es sei $\mathbf{y} = \lambda_1 \mathbf{x}_1 + \dots + \lambda_n \mathbf{x}_n$ eine Linearkombination der $\mathbf{x}_j$. Dann sind die λ_j eindeutig bestimmt.*

Beweis: Angenommen, es gebe eine zweite Menge $\mu_1, \mu_2, \dots, \mu_n$ von Koeffizienten mit

$$\mathbf{y} = \mu_1 \mathbf{x}_1 + \mu_2 \mathbf{x}_2 + \dots + \mu_n \mathbf{x}_n.$$

Subtrahiert man diese Gleichung von der obigen, so erhält man

$$\vec{0} = (\lambda_1 - \mu_1) \mathbf{x}_1 + \dots + (\lambda_n - \mu_n) \mathbf{x}_n.$$

Da aber die $\mathbf{x}_j$ als linear unabhängig vorausgesetzt worden sind, folgt, dass $\lambda_j - \mu_j = 0$, d.h. $\lambda_j = \mu_j$ für alle j , so dass nur ein Satz von Koeffizienten λ_j existiert, um $\mathbf{y}$ als Linearkombination der $\mathbf{x}_j$ darzustellen. $\square$

Beispiel 2.10 Lineare Unabhängigkeit der Einheitsvektoren Die Einheitsvektoren $\mathbf{e}_1, \dots, \mathbf{e}_n$ sind linear unabhängig. Denn

$$\vec{0} = \lambda_1 \mathbf{e}_1 + \lambda_2 \mathbf{e}_2 + \dots + \lambda_n \mathbf{e}_n. \quad (2.69)$$

Für die i -te Komponente des Nullvektors hat man nämlich

$$0 = \lambda_1 0 + \dots + \lambda_i 1 + \dots + \lambda_n 0,$$

d.h. $\lambda_i 1 = 0$, und dies ist nur möglich, wenn $\lambda_i = 0$; dies gilt für alle $i = 1, \dots, n$ (vergl. (2.19), Seite 18). $\square$

Beispiel 2.11 Lineare Unabhängigkeit und Gleichungssysteme: Gegeben seien zwei 2-dimensionale Vektoren $\mathbf{x} = (1, 2)'$ und $\mathbf{y} = (3, 4)'$. Die Frage ist, ob sie linear unabhängig sind oder nicht. Dazu betrachtet man die Gleichung $\lambda_1 \mathbf{x} + \lambda_2 \mathbf{y} = \vec{0}$; wenn sie nur gilt, wenn $\lambda_1 = \lambda_2 = 0$ ist, dann sind sie linear unabhängig. Tatsächlich repräsentiert die Gleichung ein System von zwei Gleichungen mit den Unbekannten λ_1 und λ_2 :

$$\begin{aligned} 1\lambda_1 + 3\lambda_2 &= 0 \\ 2\lambda_1 + 4\lambda_2 &= 0 \end{aligned}$$

Die erste Gleichung impliziert $\lambda_1 = -3\lambda_2$, und wenn man diesen Ausdruck für λ_1 in die zweite Gleichung einsetzt, erhält man

$$-6\lambda_2 + 4\lambda_2 = 0 \Rightarrow -2\lambda_2 = 0,$$

woraus $\lambda_2 = 0$ und damit $\lambda_1 = 0$ folgt. Die beiden Vektoren sind linear unabhängig.

Nun sei $\mathbf{y} = (2, 4)'$, und $\mathbf{x}$ sei wie oben definiert. Man erhält das Gleichungssystem

$$\begin{aligned} 1\lambda_1 + 2\lambda_2 &= 0 \\ 2\lambda_1 + 4\lambda_2 &= 0 \end{aligned}$$

Subtrahiert man die erste Gleichung von der zweiten, so erhält man die Gleichung $\lambda_1 + 2\lambda_2 = 0$, d.h. $\lambda_1 = -2\lambda_2$. Alle λ_1, λ_2 -Werte, die dieser Gleichung genügen, genügen der Vektorgleichung $\lambda_1\mathbf{x} + \lambda_2\mathbf{y} = \vec{0}$, so dass $\mathbf{x}$ und $\mathbf{y}$ linear abhängig sind. Der Grund für diesen Befund liegt natürlich darin, dass $\mathbf{y} = 2\mathbf{x}$ ist, d.h. die beiden Vektoren sind linear abhängig; sie haben dieselbe Orientierung. Dieser Befund gilt nicht nur für 2-dimensionale Vektoren: zwei Vektoren mit identischer Orientierung sind stets linear abhängig, wie in Abschnitt 2.3.2 deutlich wird.

Offenbar besteht eine Beziehung zwischen dem Begriff der linearen Unabhängigkeit bzw. Abhängigkeit und der Menge der Lösungen für das Gleichungssystem $\lambda_1\mathbf{x}_1 + \lambda_2\mathbf{x}_2 + \cdots + \lambda_n\mathbf{x}_n = \vec{0}$, – in ausgeschriebener Form bedeutet dieser Ausdruck ja

$$\begin{aligned} \lambda_1 x_{11} + \lambda_2 x_{12} + \cdots + \lambda_n x_{1n} &= 0 \\ \lambda_1 x_{21} + \lambda_2 x_{22} + \cdots + \lambda_n x_{2n} &= 0 \\ &\vdots \\ \lambda_1 x_{n1} + \lambda_2 x_{n2} + \cdots + \lambda_n x_{nn} &= 0 \end{aligned} \tag{2.70}$$

wobei die Koeffizienten $\lambda_1, \dots, \lambda_n$ die Unbekannten sind. Schreibt man $\boldsymbol{\lambda} = (\lambda_1, \lambda_2, \dots, \lambda_n)'$ und nennt $\boldsymbol{\lambda}$ den *Lösungsvektor*, so bedeutet die lineare Unabhängigkeit der $\mathbf{x}_1, \dots, \mathbf{x}_n$, dass es nur eine Lösung gibt, nämlich $\boldsymbol{\lambda} = \vec{0}$ (vergl. dazu auch Satz 2.1). Daraus folgt, dass die Existenz einer Lösung $\boldsymbol{\lambda} \neq \vec{0}$ impliziert, dass die $\mathbf{x}_1, \dots, \mathbf{x}_n$ linear abhängig sind. Ist $\boldsymbol{\lambda} \neq \vec{0}$ eine Lösung, so ist offenbar $a\boldsymbol{\lambda}$ mit $a \in \mathbb{R}$ ebenfalls eine Lösung, denn dann gilt ja

$$a\lambda_1\mathbf{x}_1 + \cdots + a\lambda_n\mathbf{x}_n = a\vec{0} = \vec{0}.$$

Sind darüber hinaus $\boldsymbol{\lambda} \neq \vec{0}$ und $\boldsymbol{\mu} \neq \vec{0}$ nichtparallele Lösungen, so ist auch jede Linearkombination $a\boldsymbol{\lambda} + b\boldsymbol{\mu}$ mit $a, b \in \mathbb{R}$ eine Lösung, wovon man sich leicht überzeugt: nach Voraussetzung muß ja $a\lambda_1\mathbf{x}_1 + \cdots + a\lambda_n\mathbf{x}_n = \vec{0}$ und $b\mu_1\mathbf{x}_1 + \cdots + b\mu_n\mathbf{x}_n = \vec{0}$ gelten und $\vec{0} + \vec{0} = \vec{0}$. Die lineare Abhängigkeit bzw. Unabhängigkeit von Vektoren steht also in einem engen Zusammenhang mit der Eindeutigkeit der Lösungen für lineare Gleichungssysteme. $\square$

2.3.2 Lineare Unabhängigkeit und Skalarprodukt

Eine notwendige, wenn auch nicht hinreichende Bedingung für die lineare Unabhängigkeit von Vektoren $\mathbf{x}_1, \dots, \mathbf{x}_p$ ist, dass sie sich paarweise hinsichtlich ihrer Orientierungen voneinander unterscheiden; dieser Sachverhalt soll jetzt etwas expliziter illustriert werden.

Es seien $\mathbf{x}$ und $\mathbf{y}$ zwei n -dimensionale Vektoren mit dem Skalarprodukt $\mathbf{x}'\mathbf{y}$. Nach (2.53) (Seite 28) gilt

$$\mathbf{x}'\mathbf{y} \leq \|\mathbf{x}\|\|\mathbf{y}\|.$$

Es gibt zwei Fälle:

1. Die $\mathbf{x}$ und $\mathbf{y}$ seien linear abhängig. Dann gilt

$$\lambda_1 \mathbf{x} + \lambda_2 \mathbf{y} = \vec{0}, \quad \lambda_1, \lambda_2 \neq 0.$$

Es folgt

$$\mathbf{y} = -\frac{\lambda_2}{\lambda_1} \mathbf{x} = \lambda \mathbf{x}, \quad \lambda = -\lambda_2/\lambda_1.$$

Dies ist aber der in (2.52) (Seite 28) betrachtete Fall, d.h. es folgt $\mathbf{x}'\mathbf{y} = \|\mathbf{x}\|\|\mathbf{y}\|$, d.h. lineare Abhängigkeit zweier Vektoren bedeutet (i) $\theta = 0$ und damit $\cos \theta = 1$, d.h. die beiden Vektoren sind parallel, und das Skalarprodukt nimmt den Maximalwert $\|\mathbf{x}\|\|\mathbf{y}\|$ an; man sagt auch, die beiden Vektoren seien *kollinear*; Anfangs- und Endpunkte der Vektoren liegen dann auf einer Geraden (vergl. Definition 2.2, Seite 11). Umgekehrt impliziert der Fall $\theta = 0$, also die Parallelität der beiden Vektoren, die lineare Abhängigkeit von $\mathbf{x}$ und $\mathbf{y}$, wie man durch Einsetzen von $\mathbf{y} = \lambda \mathbf{x}$, $\lambda \neq 0$, sofort sieht.

2. Die Vektoren $\mathbf{x}$ und $\mathbf{y}$ seien linear unabhängig. Dann gilt $\lambda_1 \mathbf{x} + \lambda_2 \mathbf{y} = \vec{0}$ nur für $\lambda_1 = \lambda_2 = 0$, d.h. es existiert kein λ mit $\mathbf{y} = \lambda \mathbf{x}$. Also sind $\mathbf{x}$ und $\mathbf{y}$ nicht parallel. Mithin folgt $\theta \neq 0$, $\cos \theta < 1$ und die beiden Vektoren haben verschiedene Orientierungen. Umgekehrt impliziert $\theta \neq 0$ und damit $|\cos \theta| < 1$ die lineare Unabhängigkeit von $\mathbf{x}$ und $\mathbf{y}$, da ja nun die Beziehung $\mathbf{y} = \lambda \mathbf{x}$ *nicht* gilt.

Es seien nun $\mathbf{x}_1, \dots, \mathbf{x}_p$ linear unabhängige Vektoren, $\mathbf{x}_j \neq \vec{0}$ für $1 \leq j \leq p$. Dann impliziert $\lambda_1 \mathbf{x}_1 + \dots + \lambda_p \mathbf{x}_p = \vec{0}$ die Beziehungen $\lambda_1 = \lambda_2 = \dots = \lambda_p = 0$. Dann können keine zwei dieser Vektoren zueinander parallel sein. Denn angenommen, es gelte $\mathbf{x}_2 = \lambda \mathbf{x}_1$ mit $\lambda \neq 0$. Dann hat man

$$\lambda_1 \mathbf{x}_1 + \lambda_2 a \mathbf{x}_1 + \lambda_3 \mathbf{x}_3 + \dots + \lambda_p \mathbf{x}_p = (\lambda_1 + a \lambda_2) \mathbf{x}_1 = \vec{0},$$

da nach Voraussetzung $\lambda_1 = \dots = \lambda_p = 0$, d.h. $\lambda_1 = -\lambda \lambda_2$, also

$$\lambda = -\frac{\lambda_1}{\lambda_2},$$

so dass λ nicht existiert, da der Ausdruck $0/0$ bekanntlich keinen Sinn macht. Die lineare Unabhängigkeit der $\mathbf{x}_1, \dots, \mathbf{x}_p$ impliziert, dass keine zwei dieser Vektoren parallel zueinander sind und damit alle p verschiedene Orientierungen haben.

Der Begriff der 'verschiedenen Orientierung' muß aber noch spezifiziert werden, um ihn mit dem der linearen Unabhängigkeit zu verknüpfen. Denn die

$\mathbf{x}_1, \dots, \mathbf{x}_p$ seien 3-dimensionale Vektoren, die alle in einer Ebene liegen. Innerhalb dieser Ebene können sie alle verschiedene Orientierungen haben, gleichwohl zeigt sich, dass sie linear abhängig sind, wie in Abschnitt 2.4.1, Beispiel 2.14 gezeigt wird.

2.3.3 Lineare Unabhängigkeit und Korrelationen

Es seien $\mathbf{x}$ und $\mathbf{y}$ irgendzwei n -dimensionale Vektoren, die einen Winkel θ einschließen. Repräsentieren die Komponenten x_i und y_i von $\mathbf{x}$ und $\mathbf{y}$ Abweichungen vom jeweiligen Mittelwert, so kann $\mathbf{x}'\mathbf{y}/(\|\mathbf{x}\|\|\mathbf{y}\|) = \cos \theta$ als Korrelationskoeffizient r_{xy} interpretiert werden. Für $-1 < r_{xy} < 1$ (oder $|r_{xy}| < 1$) sind $\mathbf{x}$ und $\mathbf{y}$ linear unabhängig. Für $b = 0$ ist $\mathbf{y} = \mathbf{e}$, also nicht aus $\mathbf{x}$ berechenbar und deswegen linear unabhängig von $\mathbf{x}$. Für den Fall $b \neq 0$ gilt $\mathbf{y} = b\mathbf{x} + \mathbf{e}$. Für den Fall $\mathbf{e} = \vec{0}$ gilt $\mathbf{y} = b\mathbf{x}$, $\mathbf{x}$ und $\mathbf{y}$ sind dann parallel und damit linear abhängig (die beiden Vektoren sind kollinear, und $r_{xy} = \pm 1$, je nachdem, ob $b > 0$ oder $b < 0$). Es sei $\mathbf{e} \neq \vec{0}$ (der übliche Fall). $\mathbf{e}$ und $\mathbf{x}$ sind aber linear unabhängig, ebenso $\mathbf{e}$ und $\mathbf{y}$, denn wären etwa $\mathbf{e}$ und $\mathbf{x}$ linear abhängig, so hätte man $\mathbf{e} = \alpha\mathbf{x}$ und $\mathbf{y} = b\mathbf{x} + \alpha\mathbf{x} = (b + \alpha)\mathbf{x} = \beta\mathbf{x}$ und es gäbe gar keinen Fehler. Also müssen $\mathbf{x}$ und $\mathbf{e}$ linear unabhängig sein; ein analoges Argument gilt für $\mathbf{e}$ und $\mathbf{y}$. Die lineare Unabhängigkeit von $\mathbf{e}$ von $\mathbf{x}$ und $\mathbf{y}$ bedeutet, dass $\mathbf{y}$ nicht aus $\mathbf{x}$ allein berechenbar ist, so dass $\mathbf{x}$ und $\mathbf{y}$ linear unabhängig sind.

Die KQ-Schätzung¹⁰ $\hat{b}$ für b ist bekanntlich

$$\hat{b} = \frac{Kov(x, y)}{s_x^2} = \frac{\mathbf{x}'\mathbf{y}}{\|\mathbf{x}\|^2}. \quad (2.71)$$

Selbst für $b = 0$ ist die Wahrscheinlichkeit, dass $\hat{b} = 0$ ist, gleich Null¹¹. Man wird irgendwie entscheiden müssen, ob $\hat{b} \neq 0$ nun in Wirklichkeit $b \neq 0$ oder $b = 0$ bedeutet. Ein Signifikanztest mag hilfreich sein, um eine Entscheidung zu treffen, liefert aber keine Gewißheit.

Es ist

$$\mathbf{y} = \hat{\mathbf{y}} + \hat{\mathbf{e}}, \quad \hat{\mathbf{y}} = \hat{b}\mathbf{x} \quad (2.72)$$

$\hat{\mathbf{e}} = \mathbf{y} - \hat{\mathbf{y}}$ ist der kleinste Fehler (im Sinne von $\hat{\mathbf{e}}'\hat{\mathbf{e}} = \min$, also der kürzeste Fehlervektor), der sich aus der auf die vorliegende Stichprobe angewandten Methode der Kleinsten Quadrate ergibt:

$$\sum_{i=1}^n (y_i - \hat{y}_i)^2 = (\mathbf{y} - \hat{\mathbf{y}})'(\mathbf{y} - \hat{\mathbf{y}}) = \|\mathbf{y} - \hat{\mathbf{y}}\|^2 = \|\hat{\mathbf{e}}\|^2. \quad (2.73)$$

¹⁰KQ-Schätzung = Kleinste Quadrate-Schätzung, s. Abschnitt 3.10.3, Seite 121.

¹¹Die Komponenten von $\mathbf{b}$ nehmen Werte auf einem Kontinuum an, und die Wahrscheinlichkeit, dass ein spezieller Wert $b \in \mathbb{R}$ auftritt, ist bei stetigen Variablen gleich Null.

Die Orthogonalität von $\hat{e}$ und $\hat{y}$: $\hat{e}$ ist nicht nur linear unabhängig von $\mathbf{x}$ und $\mathbf{y}$, sondern $\hat{e}$ und $\hat{y}$ sind orthogonal, wenn $\hat{y}$ die "Vorhersage" von $\mathbf{y}$ anhand der KQ-Schätzung $\hat{b}$ für b ist:

$$\begin{aligned}\hat{\mathbf{y}}'\hat{\mathbf{e}} &= \hat{\mathbf{y}}'(\mathbf{y} - \hat{\mathbf{y}}) = \hat{\mathbf{y}}'\mathbf{y} - \hat{\mathbf{y}}'\hat{\mathbf{y}} \\ &= \frac{\mathbf{x}'(\mathbf{x}'\mathbf{y})\mathbf{y}}{\|\mathbf{x}\|^2} - \frac{\mathbf{x}'(\mathbf{x}'\mathbf{y})^2\mathbf{x}}{\|\mathbf{x}\|^4} = \frac{(\mathbf{x}'\mathbf{y})^2}{\|\mathbf{x}\|^2} - \frac{(\mathbf{x}'\mathbf{y})^2}{\|\mathbf{x}\|^2} = 0\end{aligned}\quad (2.74)$$

Multipliziert man (2.72) mit $\mathbf{y}'$, so erhält man

$$\mathbf{y}'\mathbf{y} = \|\mathbf{y}\|^2 = \mathbf{y}'\hat{\mathbf{y}} + \mathbf{y}'\hat{\mathbf{e}}.$$

Division durch $\|\mathbf{y}\|^2$ liefert wegen $\hat{y} = \hat{b}x$ und (2.71)

$$1 = \frac{(\mathbf{x}'\mathbf{y})^2}{(\|\mathbf{x}\|\|\mathbf{y}\|)^2} + \frac{\|\hat{\mathbf{e}}\|^2}{\|\mathbf{y}\|^2}\quad (2.75)$$

bzw

$$r_{xy}^2 = 1 - \frac{\|\hat{\mathbf{e}}\|^2}{\|\mathbf{y}\|^2} \quad \text{bzw} \quad \frac{\|\hat{\mathbf{e}}\|^2}{\|\mathbf{y}\|^2} = 1 - r_{xy}^2\quad (2.76)$$

Der Determinationskoeffizient r_{xy}^2 liefert also eine Abschätzung des Anteils $\|\hat{\mathbf{e}}\|^2/\|\mathbf{y}\|^2$, dh des Anteils der Fehlervarianz an der Varianz der Variablen y . Der Wert von r_{xy}^2 erweist sich bei der Bewertung von $\hat{b}$ als nützlich.

2.3.4 Lineare Unabhängigkeit und Orthogonalität

Es gilt:

1. Sind Vektoren paarweise orthogonal, so sind sie auch linear unabhängig.
2. Sind Vektoren linear unabhängig, so sind sie *nicht* notwendig auch orthogonal.

Beweis der ersten Behauptung: Die n -dimensionalen Vektoren $\mathbf{x}_1, \dots, \mathbf{x}_p$ seien paarweise orthogonal, d.h. es gelte $\mathbf{x}'_j\mathbf{x}_k = 0$ für $j \neq k$. Dann sind die $\mathbf{x}_1, \dots, \mathbf{x}_p$ auch linear unabhängig. Denn es gelte gemäß (2.68) die Gleichung

$$\vec{0} = \lambda_1\mathbf{x}_1 + \lambda_2\mathbf{x}_2 + \dots + \lambda_p\mathbf{x}_p.$$

Um zu zeigen, dass die Vektoren linear unabhängig sind, muß man zeigen, dass die λ_k alle gleich Null sind. Man wähle einen der Vektoren, etwa $\mathbf{x}_j$, und multipliziere die Gleichung mit $\mathbf{x}'_j$:

$$\mathbf{x}'_j\vec{0} = \vec{0} = \lambda_1\mathbf{x}'_j\mathbf{x}_1 + \lambda_2\mathbf{x}'_j\mathbf{x}_2 + \dots + \lambda_p\mathbf{x}'_j\mathbf{x}_p.\quad (2.77)$$

Wegen der vorausgesetzten paarweisen Orthogonalität verschwinden auf der rechten Seite alle Skalarprodukte mit die $j \neq k$, und nur für $j = k$ ist $\mathbf{x}'_j\mathbf{x}_j = \|\mathbf{x}_j\|^2 \neq$

0 (es wird vorausgesetzt, dass keiner der Vektoren der Nullvektor ist, also $\mathbf{x}_j \neq \vec{0}$ für alle j). Dann hat man

$$0 = \lambda_j \|\mathbf{x}_j\|^2,$$

und wegen $\|\mathbf{x}_j\| \neq 0$ folgt $\lambda_j = 0$. Dies gilt für alle j , so dass alle λ_j gleich Null sind, und damit sind die $\mathbf{x}_1, \dots, \mathbf{x}_p$ linear unabhängig.

Sind die $\mathbf{x}_1, \dots, \mathbf{x}_n$ l. u., so folgt, dass $\lambda_1 = \dots = \lambda_n = 0$. (2.77) folgt auch dann, wenn keines der Skalarprodukte auf der rechten Seite verschwindet, da ja alle $\lambda_j = 0$. Also folgt aus der linearen Unabhängigkeit nicht, dass die $\mathbf{x}_j$ auch paarweise orthogonal sind.

Beispiel 2.12 Die Einheitsvektoren $\mathbf{e}_j$ sind orthogonal, da $\mathbf{e}_j' \mathbf{e}_k = 0$, $j \neq k$, denn bei $\mathbf{e}_j$ steht die 1 an der j -ten, bei $\mathbf{e}_k$ an der k -ten Stelle und die korrespondierenden Produkte sind $1 \cdot 0 = 0$ und $\mathbf{e}_k = k \cdot 0 = 0$, so dass alle Produkte des Skalarprodukts gleich Null sind. Man bemerke, dass die Orthogonalität nicht aus der linearen Unabhängigkeit der $\mathbf{e}_j$ und $\mathbf{e}_k$ folgt, sondern in diesem Fall die lineare Unabhängigkeit aus der Orthogonalität. $\square$

2.4 Vektorräume

n -dimensionale Datenvektoren sind eine Teilmenge der Menge V_n aller n -dimensionalen Vektoren; die Menge V_n bildet einen *n -dimensionalen Vektorraum*. Es zeigt sich, dass die Elemente eines Vektorraumes sich als Linearkombinationen von n linear unabhängigen, n -dimensionalen Vektoren $\mathbf{b}_1, \dots, \mathbf{b}_n$ darstellen lassen; diese Vektoren liegen nicht eindeutig fest, es kann jede Menge von n linear unabhängigen Vektoren gewählt werden. Werden nur $r < n$ linear unabhängige Vektoren gewählt, so bildet die Menge der von den $\mathbf{b}_1, \dots, \mathbf{b}_r$ durch Linearkombination erzeugten Vektoren einen Teilvektorraum. Eine Menge von linear unabhängigen Vektoren $\mathbf{b}_1, \dots, \mathbf{b}_r$, $r \leq n$ bildet eine *Basis* bzw. eine *Teilbasis* im Fall $r < n$. In der multivariaten Analyse werden oft latente Vektoren für die beobachteten Vektoren gesucht; die $\mathbf{b}_j$ können als solche latenten Vektoren betrachtet werden. Die Bedeutung des Begriffs des Vektorraums wird im Folgenden elaboriert.

2.4.1 Der Begriff des Vektorraums

Gegeben sei eine Menge V_n von n -dimensionalen Vektoren. Für beliebige $\mathbf{x}, \mathbf{y} \in V_n$ und $\lambda, \mu \in \mathbb{R}$ gelte $\mathbf{z} = \lambda \mathbf{x} + \mu \mathbf{y} \in V_n$, d.h. eine beliebige Linearkombination der $\mathbf{x}, \mathbf{y}$ sei ebenfalls ein Vektor aus V_n ; man sagt, V_n sei *abgeschlossen* gegenüber der Bildung beliebiger Linearkombinationen. Dann heißt die Menge V_n *Vektorraum*, insbesondere ein *n -dimensionaler Vektorraum*.

Offenbar ist nicht jede Menge von Vektoren ein Vektorraum. So sei W eine Menge n -dimensionaler Vektoren, die alle die Länge 1 haben. Es seien $\mathbf{x}$ und $\mathbf{y}$

irgendzwei nicht-orthogonale Vektoren aus W . Dann ist

$$\|\mathbf{x} + \mathbf{y}\|^2 = (\mathbf{x} + \mathbf{y})'(\mathbf{x} + \mathbf{y}) = \|\mathbf{x}\|^2 + \|\mathbf{y}\|^2 + \mathbf{x}'\mathbf{y} = 2(1 + \mathbf{x}'\mathbf{y}) \neq 1,$$

d.h. die Summe zweier Vektoren hat nicht dieselbe Länge wie die Vektoren selbst. W ist deshalb kein Vektorraum.

Wie der Vektorbegriff ist auch der Begriff des Vektorraums sehr viel allgemeiner, als er hier zunächst eingeführt wird, in Fischer (1997) findet man eine sehr allgemeine Einführung des Begriffs des Vektorraums. Für die Zwecke der gewöhnlichen multivariaten Statistik ist der hier vorgestellte Begriff allerdings hinreichend; im Anhang, Abschnitt 8.3, findet man die allgemeine Definition.

Definition 2.8 *Es sei V_n eine Menge von n -dimensionalen Vektoren. Es gelte*

1. *Für $\mathbf{x}, \mathbf{y} \in V_n$ ist auch $\mathbf{x} + \mathbf{y} \in V_n$,*
2. *$\lambda \in \mathbb{R}, \mathbf{x} \in V_n \Rightarrow \lambda \mathbf{x} \in V_n$,*
3. *$\mathbf{x}, \mathbf{y} \in V_n, \mathbf{x} + \mathbf{y} = \mathbf{y} + \mathbf{x}$,*
4. *$\lambda \in \mathbb{R}, \mathbf{x}, \mathbf{y} \in V_n, \lambda(\mathbf{x} + \mathbf{y}) = \lambda \mathbf{x} + \lambda \mathbf{y}$*
5. *$\lambda, \mu \in \mathbb{R}, \lambda(\mu \mathbf{x}) = (\lambda \mu) \mathbf{x} \in V_n$ für $\mathbf{x} \in V_n$,*
6. *$1 \cdot \mathbf{x} = \mathbf{x}$ für $\mathbf{x} \in V_n$.*

Dann ist V_n ein n -dimensionaler Vektorraum.

Die Definition scheint einige Selbstverständlichkeiten zu enthalten; man sieht sofort, dass die Bedingungen 3. bis 6. für Vektoren mit reellen Komponenten alle erfüllt sind. Bei den in Abschnitt 8.3 betrachteten Verallgemeinerungen ist das nicht notwendig der Fall.

Ein ebenfalls wichtiger Begriff ist der des Untervektorraums:

Definition 2.9 *Es sei $U \subset V$ eine Teilmenge des Vektorraums V . $\mathbb{K}$. U ist genau dann ein Teilvektorraum oder auch Unterraum von V , wenn (i) $U \neq \emptyset$ (d.h. keine leere Menge ist), und (ii) die Bedingungen*

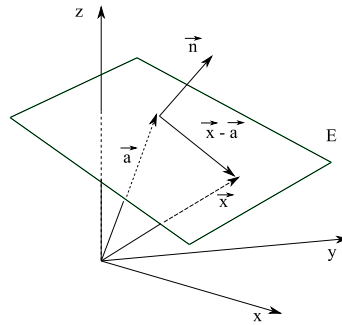
1. *$\mathbf{x}, \mathbf{y} \in U \Rightarrow \mathbf{x} + \mathbf{y} \in U$,*
2. *$c \in \mathbb{K}, \mathbf{x} \in U \Rightarrow c\mathbf{x} \in U$.*

erfüllt sind.

Ein Teilraum eines Vektorraums ist also eine Teilmenge von Vektoren derart, dass jede Linearkombination von Vektoren aus dieser Teilmenge wieder ein Element aus dieser Teilmenge ist.

Beispiel 2.13 Geraden im n -dimensionalen Vektorraum: Ein einfaches Beispiel für einen Teilraum ist eine Gerade. Ist $\mathbf{x}$ ein n -dimensionaler Vektor und $a \in \mathbb{R}$, so ist $\mathbf{u} = a\mathbf{x}$ parallel zu $\mathbf{x}$ bzw. $\mathbf{u}$ liegt auf derselben Geraden

Abbildung 5: Ebene im dreidimensionalen Raum



wie $\mathbf{x}$. Gilt $\mathbf{v} = b\mathbf{x}$, so ist eine beliebige Linearkombination von $\mathbf{u}$ und $\mathbf{v}$ durch $a_1a\mathbf{x} + a_2b\mathbf{x} = (a_1a + a_2b)\mathbf{x}$ gegeben. Sie liegt damit wieder auf der durch $\mathbf{x}$ definierten Geraden. $\square$

Beispiel 2.14 Ebenen im n -dimensionalen Vektorraum: Es sei V_n der Vektorraum aller n -dimensionalen Vektoren, und $\mathbf{x}_1$ und $\mathbf{x}_2$ zwei nicht-parallele Elemente aus V_n . Sie definieren eine Ebene in V_n . Denn es sei $\mathbf{n}$ ein Vektor aus V_n , der orthogonal zu $\mathbf{x}_1$ und $\mathbf{x}_2$ ist. Dann ist jede Linearkombination $\mathbf{x}$ von $\mathbf{x}_1$ und $\mathbf{x}_2$ ebenfalls orthogonal zu $\mathbf{n}$:

$$\mathbf{n}'\mathbf{x} = a_1\mathbf{n}'\mathbf{x}_1 + a_2\mathbf{n}'\mathbf{x}_2 = 0,$$

da ja nach Voraussetzung $\mathbf{n}'\mathbf{x}_1 = \mathbf{n}'\mathbf{x}_2 = 0$. $\mathbf{n}$ heißt *Normalenvektor*; seine Länge ist nicht relevant; im Prinzip ist es deshalb möglich, ihn als normiert zu betrachten, d.h. $\|\mathbf{n}\| = 1$ zu setzen, aber dies ist keine Notwendigkeit. Für einen Vektor $\mathbf{x}$ aus der Ebene gilt also

$$\mathbf{n}'\mathbf{x} = n_1x_1 + n_2x_2 + \cdots + n_nx_n = 0; \quad (2.78)$$

alle Vektoren $\mathbf{x}$, die dieser Gleichung genügen, sind Element der Ebene. (2.78) heißt auch *Ebenengleichung*, wobei die Ebene durch den Nullpunkt des Koordinatensystems geht, Abbildung 5 illustriert den allgemeinen Fall einer Ebene im 3-dimensionalen Raum, für $\mathbf{a} = \vec{0}$ erhält man den Fall (2.78). Den allgemeinen Fall einer Ebene erhält man, indem man eine durch (2.78) beschriebene Ebene durch den Nullpunkt durch einen Stützvektor $\mathbf{a}$ vom Nullpunkt trennt. Die Koordinatengleichung

$$n_1x_1 + n_2x_2 + \cdots + n_nx_n = a \quad (2.79)$$

beschreibt für $a \neq 0$ den allgemeinen Fall einer nicht durch den Nullpunkt gehenden Ebene. Eine andere Schreibweise für eine Ebene ist

$$\mathcal{E} = \{(x_1, \dots, x_n) \in \mathbb{R}^n \mid \mathbf{n}'\mathbf{x} = \sum_{i=1}^n n_i x_i = a\}, \quad (2.80)$$

wobei $\mathcal{E}$ die Menge der Punkte ist, die in der Ebene liegen, die durch $\mathbf{n}'\mathbf{x} = a$ liegen. Der Abstand der Ebene $\mathcal{E}$ vom Nullpunkt ist durch

$$\text{Abstand} = \frac{|a|}{\|\mathbf{n}\|}$$

gegeben.

Sind $\mathbf{y}_1$ und $\mathbf{y}_2$ irgendzwei Linearkombinationen von $\mathbf{x}_1, \mathbf{x}_2$ und damit Elemente der Ebene, so sind auch die Linearkombinationen von $\mathbf{y}_1, \mathbf{y}_2$ wieder Elemente der Ebene. Denn sicherlich sind insbesondere

$$\mathbf{y}_1 = a_1\mathbf{x}_1 + a_2\mathbf{x}_2, \quad \mathbf{y}_2 = b_1\mathbf{x}_1 + b_2\mathbf{x}_2$$

Elemente von V für $a_i, b_i \in \mathbb{R}$, $i = 1, 2$, da V ja *alle* Linearkombinationen von $\mathbf{x}_1$ und $\mathbf{x}_2$ enthält. Dann folgt

$$\mathbf{y} = c_1\mathbf{y}_1 + c_2\mathbf{y}_2 = c_1(a_1\mathbf{x}_1 + a_2\mathbf{x}_2) + c_2(b_1\mathbf{x}_1 + b_2\mathbf{x}_2),$$

woraus sich

$$\mathbf{y} = (c_1a_1 + c_2b_1)\mathbf{x}_1 + (c_1a_2 + c_2b_2)\mathbf{x}_2$$

ergibt, d.h. Linearkombinationen von Linearkombinationen von $\mathbf{x}_1$ und $\mathbf{x}_2$ sind ebenfalls Linearkombinationen von $\mathbf{x}_1$ und $\mathbf{x}_2$, also Elemente von V .

Die Bedingung $\mathbf{n}'\mathbf{x} = 0$ definiert die Restriktionen, denen die Komponenten von $\mathbf{x}$ genügen müssen, damit der Vektor Element der Ebene ist. Dies wird für den Fall $n = 3$ illustriert. Es sei $\mathbf{n} = (1, 2, 3)'$. Dann ist

$$\mathbf{n}'\mathbf{x} = 1x_1 + 2x_2 + 3x_3.$$

Gilt zum Beispiel $x_2 = 2, x_3 = 1$, so folgt $x_1 = -7$. $\mathbf{y} = (y_1, y_2, y_3)'$ sei ein weiterer Vektor aus der Ebene, so dass

$$y_1 = -(2y_2 + 3y_3).$$

Für $y_2 = y_3 = 2$ erhält man $y_1 = -10$. $\mathbf{z}$ sei eine Linearkombination von $\mathbf{x}$ und $\mathbf{y}$:

$$\mathbf{z} = \mathbf{x} + \mathbf{y} = (-17, 4, 3)'$$

Dann ist

$$\mathbf{n}'\mathbf{z} = -17 + 8 + 9 = 0,$$

d.h. $\mathbf{z}$ ist Element der Ebene. Oder $\mathbf{z} = a_1\mathbf{x} + a_2\mathbf{y}$ mit $a_1 = .5, a_2 = 1.5$:

$$\mathbf{z} = .5 \begin{pmatrix} -7 \\ 2 \\ 1 \end{pmatrix} + 1.5 \begin{pmatrix} -10 \\ 2 \\ 2 \end{pmatrix} = \begin{pmatrix} -18.5 \\ 4 \\ 3.5 \end{pmatrix}$$

und

$$\mathbf{n}'\mathbf{z} = -18.5 + 8 + 10.5 = 0.$$

$\mathbf{z}$ ist demnach wieder Element der Ebene.

Natürlich läßt sich für zwei gegebene, nicht-parallele Vektoren $\mathbf{x}_1$ und $\mathbf{x}_2$ der zugehörige Normalenvektor $\mathbf{n}$ bestimmen, – bis auf einen Faktor, d.h. bis auf die Länge. Wie oben schon angemerkt wurde, ist die ist aber nicht relevant, – relevant ist nur die Orthogonalität. $\square$

Die durch (2.80) allgemein definierte Ebene wird durch nur zwei beliebig gewählte, aber nicht-parallele Vektoren erzeugt, d.h. alle Vektoren der Ebene sind Linearkombinationen dieser zwei Vektoren. $\mathcal{E}$ definiert damit einen *2-dimensionalen Teilraum* des V_n , – auch, wenn die Vektoren n -dimensional sind und $n > 2$. Der im Folgenden besprochene Fall der Hyperebene ist davon zu unterscheiden.

Hyperebenen: Gegeben sei der $\mathbb{R}^3$ und zwei linear unabhängige, 3-dimensionale Vektoren; alle Linearkombinationen der Form $\mathbf{x} = a_1\mathbf{x}_1 + a_2\mathbf{x}_2$, $a_1, a_2 \in \mathbb{R}$, definieren Vektoren, die in einer Ebene des $\mathbb{R}^3$ liegen. Um *alle* 3-dimensionalen Vektoren zu erzeugen benötigt man drei linear unabhängige Vektoren. Es liegt nun nahe, den Begriff der Ebene zu verallgemeinern: gegeben sei der $\mathbb{R}^n$. Es werden Linearkombinationen von $n - 1$ linear unabhängigen Vektoren betrachtet:

$$\mathbf{x} = a_1\mathbf{x}_1 + \cdots + a_{n-1}\mathbf{x}_{n-1}. \quad (2.81)$$

Diese Gleichung definiert eine *Hyperebene*. Dies ist aber nicht die allgemeinste Definition von Hyperebenen. Denn die Ebene vom eben beschriebenen Typ kann parallel verschoben werden:

Definition 2.10 *Es sei*

$$\mathcal{H} = \{\mathbf{x} \in \mathbb{R}^n \mid \mathbf{s} + a_1\mathbf{x}_1 + \cdots + a_{n-1}\mathbf{x}_{n-1}\} \quad (2.82)$$

wobei $a_1, \dots, a_{n-1} \in \mathbb{R}$ reelle Zahlen sind und $\mathbf{s}$ ein Stützvektor ist. Dann heißt $\mathcal{H}$ eine Hyperebene im $\mathbb{R}^n$.

Es ist

$$\mathbf{x} - \mathbf{s} = a_1\mathbf{x}_1 + \cdots + a_{n-1}\mathbf{x}_{n-1},$$

und es sei $\mathbf{n}$ ein Normalenvektor, so dass

$$\mathbf{n}'(\mathbf{x} - \mathbf{s}) = 0.$$

Dementsprechend ist eine alternative Definition einer Hyperebene durch

$$\mathcal{H} = \{\mathbf{x} \in \mathbb{R}^n \mid \mathbf{n}'(\vec{x} - \mathbf{s}) = 0\}. \quad (2.83)$$

gegeben. Eine weitere Möglichkeit, Hyperebenen zu charakterisieren, besteht darin, die Menge der Vektoren $\mathbf{x}$ zu spezifizieren, für das Skalarprodukt $\mathbf{x}'\mathbf{c} = \langle \mathbf{x}, \mathbf{c} \rangle \in \mathbb{R}$ eine Konstante und $\mathbf{c}$ ein speziell gewählter Vektor $\mathbf{c} \in \mathbb{R}^n$ ist. Es sei wieder

$$\mathbf{x} = a_1\mathbf{x}_1 + \cdots + a_{n-1}\mathbf{x}_{n-1} + \mathbf{s}.$$

Für einen bestimmten Vektor $\mathbf{c} \in \mathbb{R}^n$ gilt dann

$$\mathbf{x}'\mathbf{c} = a_1\mathbf{x}'_1\mathbf{c} + \cdots + a_{n-1}\mathbf{x}'_{n-1}\mathbf{c} + \mathbf{s}'\mathbf{c}.$$

Für $\mathbf{x}_1, \dots, \mathbf{x}_{n-1}$ und $\mathbf{c}$ fix sind die

$$\alpha_j = \mathbf{x}'_j\mathbf{c}$$

fixe Konstanten und die $\mathbf{x}$ mit $\mathbf{x}'\mathbf{c} = d$, $d \in \mathbb{R}$, definieren offenbar eine Hyperebene

$$\mathcal{H} = \{\mathbf{x} | \mathbf{x}'\mathbf{c} = x_1c_1 + \cdots + x_nc_n = a_1\alpha_1 + \cdots + a_n\alpha_n = d\}. \quad (2.84)$$

Die $\mathbf{x}$ ergeben sich, wie schon in (2.81) durch Variation der a_j .

Satz 2.2 *Es sei $\mathcal{H}$ eine Hyperebene mit dem Parametervektor $\mathbf{a} = (a_1, a_2, \dots, a_{n-1})'$. Dann steht $\mathbf{a}$ senkrecht auf $\mathcal{H}$, d.h. $\mathbf{a}$ hat die Orientierung des Normalenvektors $\mathbf{n}$ der Hyperebene.*

Beweis: Es seien $\mathbf{x}_A$ und $\mathbf{x}_B$ zwei Punkte auf der Ebene; sie werden hier als Endpunkte der Vektoren $\mathbf{x}_A$ und $\mathbf{x}_B$ betrachtet. Dann liegt der Vektor $\mathbf{x}_A - \mathbf{x}_B$ in $\mathcal{H}$: Es gilt ja

$$\mathbf{x}'_A\mathbf{a} - \mathbf{s} = \mathbf{x}'_B\mathbf{a} - \mathbf{s},$$

d.h.

$$(\mathbf{x}_A - \mathbf{x}_B)' \mathbf{a} = \vec{0},$$

womit die Behauptung schon bewiesen ist. $\square$

Ein Beispiel für eine Teilmenge von Vektoren, die *kein* Teilraum ist, ist

Beispiel 2.15 Die Menge aller n -dimensionalen Vektoren mit einer bestimmten Länge, etwa $\|\mathbf{x}\| = a \in \mathbb{R}$, ist kein Teilraum des V_n , denn dazu müßte auch jede Linearkombination $\lambda\mathbf{x}_1 + \mu\mathbf{x}_2$ dieser Vektoren wieder die Länge a haben, $0 \neq \lambda, \mu \in \mathbb{R}$ beliebig. Es ist aber

$$\|\lambda\mathbf{x}_1 + \mu\mathbf{x}_2\| = \left(\sum_{i=1}^n (\lambda x_{i1} + \mu x_{i2})^2 \right)^{1/2} = (\lambda^2 \|\mathbf{x}_1\|^2 + \mu^2 \|\mathbf{x}_2\|^2 + 2\lambda\mu\mathbf{x}'_1\mathbf{x}_2)^{1/2} \neq a,$$

bis auf spezielle Vektoren $\mathbf{x}_1, \mathbf{x}_2$ und speziell gewählte Werte von λ und μ , etwa für $\mathbf{x}'_1\mathbf{x}_2 = 0$ und $(\lambda^2 + \mu^2)^{1/2} = 1$. $\square$

Satz 2.3 *Jeder Vektorraum V besitzt genau einen Nullvektor $\vec{0}$. $\vec{0}$ liegt in allen Unter- oder Teilräumen von V , und für $\mathbf{v} \in V$ und $\lambda \in \mathbb{R}$ gilt $\lambda\mathbf{v} = \vec{0}$ genau dann, wenn $\lambda = 0$ oder $\mathbf{v} = \vec{0}$.*

Beweis: Nach der Bedingung K1 für Vektorräume existiert für V ein neutrales Element $0 = \vec{0}$, so dass $\vec{0} + \mathbf{v} = \mathbf{v}$ für alle $\mathbf{v} \in V$. Gilt also für $\mathbf{w} \neq \mathbf{v} \in V$ auch $\vec{0} + \mathbf{w} = \mathbf{w} + \vec{0} = \mathbf{w}$, so folgt die Eindeutigkeit von $\vec{0}$. Ist U ein Unterraum von V , so ist $U \neq \emptyset$ und für jeden Vektor $\mathbf{u} \in U$ existiert ein $-\mathbf{u} = (-1)\mathbf{u}$ und es gilt $\mathbf{u} - \mathbf{u} = \mathbf{u} + (-\mathbf{u}) = \vec{0}$.

Weiter folgt $0 \cdot \mathbf{v} = \vec{0}$ für alle $\mathbf{v} \in V$, denn $0 \cdot \mathbf{v} + 0 \cdot \mathbf{v} = 0 \cdot \mathbf{v}$. Subtrahiert man $0 \cdot \mathbf{v}$ auf beiden Seiten dieser Gleichung, so erhält man $0 \cdot \mathbf{v} = \vec{0}$. Nach (iii) der Definition 8.4 gilt weiter

$$\lambda \cdot \vec{0} + \lambda \cdot \vec{0} = \lambda(\vec{0} + \vec{0}) = \lambda \cdot \vec{0},$$

so dass $\lambda \vec{0} = \vec{0}$. Umgekehrt sei $\lambda \mathbf{v} = \vec{0}$ und $\lambda \neq 0$. Dann existiert $\lambda^{-1} \in K$ und es gilt $\lambda \mathbf{v} \cdot 1 = (\lambda^{-1} \lambda) \mathbf{v} = \lambda^{-1}(\lambda \mathbf{v}) = \lambda^{-1} \vec{0} = \vec{0}$. $\square$

Gegeben sei ein 3-dimensionaler Vektorraum $\mathbb{R}^3$ und es werden zwei verschiedene Ebenen in diesem Raum betrachtet. Wenn die Ebenen nicht parallel sind, schneiden sie sich irgendwo; die Schnittmenge ist eine Gerade, – und eine Gerade ist wieder ein Unterraum der $\mathbb{R}^3$. Formal entspricht diese Gerade dem Durchschnitt der beiden Mengen von Vektoren, die in der einen bzw. der anderen Ebene liegen, – man benutzt für den Durchschnitt das Zeichen $\cap$. Man kann auch die Vereinigung dieser beiden Mengen betrachten; für die Vereinigung zweier Mengen benutzt man das Zeichen $\cup$. Die Frage ist, ob die Vereinigung zweier Unterräume ebenfalls wieder ein Vektorraum ist. Auskunft darüber gibt der folgende allgemeine

Satz 2.4 *Es sei V ein Vektorraum und U und W seien Teilräume von V . Dann gilt*

- (a) *$U \cap W$ ist ebenfalls ein Unterraum von V ,*
- (b) *$U \cup W$ ist genau dann ein Unterraum von V , wenn $U \subseteq W$ oder $W \subseteq U$ gilt.*

Beweis: (a): Es sei $\mathbf{v}_1, \mathbf{v}_2 \in U \cap W$. Dann sind $\mathbf{v}_1$ und $\mathbf{v}_2$ Elemente von U , und da U ein Unterraum von V ist, ist auch jede Linearkombination $a\mathbf{v}_1 + b\mathbf{v}_2$ in U . Aber $\mathbf{v}_1$ und $\mathbf{v}_2$ sind auch Elemente von W , und da W ebenfalls ein Unterraum von V ist, ist jede Linearkombination von $\mathbf{v}_1$ und $\mathbf{v}_2$ auch Element von W , d.h. die Linearkombinationen sind Elemente von $U \cap W$.

(b): Ist $U \subseteq W$ oder $W \subseteq U$, so ist $U \cup W$ ein Unterraum von V . Für $U \subseteq W$ folgt $U \cup W = W$ und aus $W \subseteq U$ folgt $U \cup W = U$, und da U und W schon Teilräume sind folgt, dass in diesen beiden Fällen auch $U \cup W$ ein Teilraum von V ist.

Nun gelte keine dieser beiden Bedingungen $U \subseteq W$ oder $W \subseteq U$. $U \cup W$ sei aber ein Unterraum von V . Wegen $U \not\subseteq W$ folgt, dass ein $\mathbf{u} \in U - W$ existiert¹²

¹²Gemeint ist die Menge U , aus der die Elemente von W entfernt sind; eine andere Schreibweise ist $U \setminus W$.

und ebenso ein Vektor $\mathbf{w} \in W - U$, und $\mathbf{u} + \mathbf{w} \in U \cup W$, d.h. $\mathbf{u} + \mathbf{w} \in U$ oder $\mathbf{u} + \mathbf{w} \in W$. Es sei nun $\mathbf{u} + \mathbf{w} \in U$. Dann muß wegen der Unterraumeigenschaft auch die Linearkombination $(\mathbf{u} + \mathbf{w}) + \lambda \mathbf{u} \in U$ gelten. Insbesondere sei $\lambda = -1$ soll also auch $(\mathbf{u} + \mathbf{w}) - \mathbf{u} = \mathbf{w} \in U$ sein. Das kann aber nicht sein, weil ja $W \subseteq U$ ja gerade nicht gelten soll. Also folgt, dass die Behauptung, $U \cup W$ sei ein Unterraum von V nicht gelten kann, zumal auch die Annahme $\mathbf{u} + \mathbf{w} \in W$ in derselben Weise zu einem Widerspruch führt. Es muß also mindestens eine der Bedingungen $U \subseteq W$ oder $W \subseteq U$ erfüllt sein, damit $U \cup W$ einen Unterraum bilden. $\square$

Anmerkung: Teil (a) des Satzes gilt nicht nur für zwei Unterräume, sondern für beliebig viele: sind die U_i , $i = 1, 2, \dots$ Unterräume von V , so ist auch $\cap U_i$ ein Unterraum. Der Beweis ist analog zu dem für (a) in Satz 2.4. $\square$

Beispiel 2.16 Beispiele: Gegeben seien zwei Ebenen mit verschiedenen Orientierungen in einem 3-dimensionalen Raum. Der Durchschnitt der Ebenen ist eine Gerade, – ein 1-dimensionaler Teilraum. Umgekehrt bilden zwei Gerade U und W im $\mathbb{R}^3$ mit verschiedener Orientierung noch keinen Teilraum, weil Linearkombinationen von Vektoren aus U einerseits und W andererseits nicht notwendig wieder in einem dieser beiden Teilräume liegen. $\square$

Definition 2.11 Es sei W ein Vektorraum und U, V seien Teilräume aus W . Dann heißt

$$S = U + V = \{\mathbf{u}, \mathbf{v} | \mathbf{u} \in U, \mathbf{v} \in V\} \quad (2.85)$$

die Summe der Vektorräume U und V .

Dann folgt der

Satz 2.5 Es sei W ein Vektorraum und U, V seien Teilräume aus W . Die Summe $U + V$ ist der kleinste Teilraum von W , der U und V enthält.

Beweis: Es seien $\mathbf{a}_1, \mathbf{a}_2 \in S$. Dann existieren $\mathbf{u}_1, \mathbf{u}_2 \in U$, $\mathbf{v}_1, \mathbf{v}_2 \in V$ mit

$$\mathbf{a}_1 = \mathbf{u}_1 + \mathbf{v}_1, \quad \mathbf{a}_2 = \mathbf{u}_2 + \mathbf{v}_2$$

mit $\mathbf{a}_1 + \mathbf{a}_2 \in S$ und für $\lambda \in \mathbb{K}$ ($\lambda \in \mathbb{R}$), $\lambda \mathbf{a} \in S$ für $\mathbf{a} \in S$, – dies folgt aus den Teilraumeigenschaften von U und V . Schließlich folgt aus $\vec{0} \in V$, $\vec{0} \in W$ auch $\vec{0} + \vec{0} = \vec{0} \in S$. Damit ist S ein (Teil-)Vektorraum. Weiter sei $Z \subseteq W$ mit $U \subseteq Z, V \subseteq Z$. Für alle $\mathbf{u} \in U$, $\mathbf{v} \in V$ und $\mathbf{z} = \mathbf{u} + \mathbf{v}$ folgt $\mathbf{z} \in Z$, d.h. aber $V + W \subseteq Z$, womit der Satz bewiesen ist. $\square$

Es zeigt sich, dass sich alle Vektoren des Vektorraums als Linearkombinationen von bestimmten Teilmengen von Vektoren erzeugen lassen. Die folgende Definition dient der Entwicklung dieses Sachverhalts.

Definition 2.12 Es sei $M = \{\mathbf{x}_1, \dots, \mathbf{x}_m\} \subseteq V$ eine Teilmenge von Vektoren eines Vektorraums V . Dann heißt

$$\mathcal{L}(M) = \{\mathbf{x} | \mathbf{x} = a_1 \mathbf{x}_1 + \dots + a_m \mathbf{x}_m; a_j \in \mathbb{K}, j = 1, \dots, m\} \quad (2.86)$$

die lineare Hülle von M .

$\mathcal{L}$ ist also die Menge der Linearkombinationen, die aus den Elementen von M erzeugt werden können. Man sagt auch, $\mathcal{L}$ wird durch die Vektoren von M *aufgespannt* und schreibt statt $\mathcal{L}(M)$ auch $\text{Spann}(M)$ oder $\text{span}(M)$; diese Schreibweise ist aus dem Englischen – to span heißt aufspannen – adaptiert.

Es sei $M \subseteq V$; M muß kein Teilraum sein, aber es gilt der

Satz 2.6 Es sei V ein Vektorraum und $M = \{\mathbf{x}_1, \dots, \mathbf{x}_m\}$ sei eine Teilmenge von Vektoren aus V . Dann ist $\mathcal{L}(M)$ ein (Teil-)Vektorraum von V .

Beweis: Es seien $\mathbf{x}$ und $\mathbf{y}$ Linearkombinationen von Vektoren aus M ,

$$\mathbf{x} = \sum_{j=1}^m a_j \mathbf{x}_j, \quad \mathbf{y} = \sum_{j=1}^m b_j \mathbf{x}_j.$$

Dann folgt

$$\lambda \mathbf{x} + \mu \mathbf{y} = \lambda \sum_{j=1}^m a_j \mathbf{x}_j + \mu \sum_{j=1}^m b_j \mathbf{x}_j = \sum_{j=1}^m c_j \mathbf{x}_j, \quad c_j = \lambda a_j + \mu b_j$$

d.h. $\lambda \mathbf{x} + \mu \mathbf{y}$ ist ebenfalls eine Linearkombination der $\mathbf{x}_j$ und damit Element von $\mathcal{L}(M)$. $\square$

Die lineare Hülle $\mathcal{L}$ einer Teilmenge M von Vektoren eines Vektorraums hat eine interessante Eigenschaft, die im folgenden Satz formuliert wird:

Satz 2.7 Es sei M eine beliebige Teilmenge eines Vektorraums V . Dann ist $U = \mathcal{L}(M)$ die kleinste Teilmenge von V , die M enthält; U ist eindeutig bestimmt.

Beweis: Aus $M \subseteq V$ folgt $U = \mathcal{L}(M) \subseteq V$, denn für jeden Vektor $\mathbf{v} \in M$ folgt $\mathbf{v} \in \mathcal{L}(M)$ und $\mathbf{v} \in V$ aufgrund der Definition von $\mathcal{L}(M)$. Es sei weiter $W \subseteq V$ derart, dass $M \subseteq W$. Für jeden Vektor $\mathbf{v} \in M$ folgt dann $\mathbf{v} \in W$ und $U \subseteq W$, d.h. U ist in *allen* Teilräumen W enthalten, die M als Teilmenge enthalten, so dass U der kleinste Teilraum ist, der M enthält. U ist eindeutig bestimmt. Denn angenommen, es gäbe einen weiteren solchen Teilraum $U' \neq U$. Aus der Definition von U folgt aber $U \subseteq U'$, und aus dem gleichen Grund folgt auch $U' \subseteq U$. Aus $U' \subseteq U$ und $U \subseteq U'$ folgt dann $U = U'$, d.h. U ist die einzige kleinste Teilmenge, die M enthält. $\square$

Der Begriff der linearen Unabhängigkeit ist zwar schon für den Spezialfall von Vektoren als n -tupeln $(x_1, \dots, x_n)'$ von Zahlen definiert worden. Hier wird er noch einmal für den allgemeinen Fall von Vektoren als Elementen von Vektorräumen charakterisiert.

Definition 2.13 Es sei V_0 eine Teilmenge des Vektorraums V . V_0 heißt linear unabhängige Teilmenge von V , wenn für jedes $\mathbf{v} \in V_0$ der von $V_0 - \{\mathbf{v}\}$ erzeugte Teilraum eine echte Teilmenge des von V_0 erzeugten Teilraums ist, so dass für alle $\mathbf{v} \in V_0$ für den von $V_0 - \{\mathbf{v}\}$ aufgespannten Raum

$$\mathcal{L}(V_0 - \{\mathbf{v}\}) \neq \mathcal{L}(V_0)$$

gilt. V_0 heißt linear abhängig, wenn V_0 nicht linear unabhängig ist.

Anmerkung: Die Redeweise vom "erzeugten Teilraum" bedeutet, dass die Menge der Linearkombinationen der Vektoren aus in diesem Fall V_0 oder $V_0 - \{\mathbf{v}\}$, also $\mathcal{L}(V_0)$ oder $\mathcal{L}(V_0 - \{\mathbf{v}\})$ betrachtet wird. $\square$

Zur Erinnerung: $\mathcal{L}(V_0)$ ist der von den Vektoren in V_0 aufgespannte Vektorraum; dieser Raum kann der Vektorraum V sein oder aber nur eine Teilmenge von V , – dies hängt von den Vektoren in V_0 ab. Entfernt man einen Vektor – etwa $\mathbf{v}$ – aus V_0 , so kann der von $V_0 - \{\mathbf{v}\}$ aufgespannte Vektorraum gleich $\mathcal{L}(V_0)$ sein oder nicht. Ist er es *nicht*, so bedeutet dies, dass der entnommene Vektor $\mathbf{v}$ gewissermaßen Information enthält, die in den anderen Vektoren nicht enthalten ist und die nun fehlt. Deshalb kann $\mathbf{v}$ eben nicht von den anderen Vektoren in V_0 "vorhergesagt", d.h. als Linearkombination berechnet werden, $\mathbf{v}$ ist in diesem Sinne unabhängig von den übrigen Vektoren in V_0 .

Definition 2.14 Es sei V_0 eine Teilmenge von V . Mit $V_0 \cup \{\mathbf{v}\}$ ist die Menge gemeint, die entsteht, wenn man ihr einen Vektor $\mathbf{v} \in V$ hinzufügt. V_0 heißt eine maximale linear unabhängige (l.u.) Teilmenge einer Teilmenge V'_0 von V , wenn V_0 linear unabhängig ist, aber $V_0 \cup \{\mathbf{v}\}$ für jeden Vektor $\mathbf{v} \in V'_0$ linear abhängig ist, wenn $\mathbf{v}$ nicht zu V_0 gehört.

Satz 2.8 Ist die Menge $V_0 \subset V$ linear unabhängig und $V_0 \cup \{\mathbf{v}\}$, $\mathbf{v} \in V$, linear abhängig, so ist $\mathbf{v} \in \mathcal{L}(V_0)$. Ist V_0 eine maximale linear unabhängige Teilmenge einer Teilmenge V'_0 von V , so ist $V'_0 \subseteq \mathcal{L}(V_0)$.

Beweis: Der erste Teil des Satzes ist intuitiv sofort klar: wenn $V_0 \cup \{\mathbf{v}\}$ linear abhängig ist, so heißt dies ja, dass $\mathbf{v}$ als Linearkombination der Vektoren in V_0 dargestellt werden kann, und dieser Sachverhalt wird gerade durch die Aussage, dass $\mathbf{v}$ in der linearen Hülle von V_0 enthalten sein muß ausgedrückt. Wenn $V_0 \subset V'_0$ maximal linear unabhängig ist, so ist $V_0 \cup \{\mathbf{v}\}$ mit $\mathbf{v} \in V'_0$, $\mathbf{v} \notin V_0$ linear abhängig und folglich muß $\mathbf{v}$ eine Linearkombination der Vektoren in V_0 sein, und damit ist $\mathbf{v} \in \mathcal{L}(V_0)$, also in der linearen Hülle von V_0 , und dies bedeutet $V'_0 \subseteq \mathcal{L}(V_0)$. $\square$

2.4.2 Basen von Vektorräumen und Teilvektorräumen

Der Begriff der linearen Hülle einer Menge M von Vektoren legt nahe, dass alle Vektoren einer Menge als Linearkombinationen bestimmter Vektoren $\mathbf{b}_1, \dots, \mathbf{b}_n$

dargestellt bzw. erzeugt werden können. Dies führt zu der folgenden Begriffsbildung:

Definition 2.15 *Es sei V ein Vektorraum. Die Teilmenge $\mathcal{B} = \{\mathbf{b}_1, \dots, \mathbf{b}_n\}$ von V heißt Basis von V , wenn gilt*

- (i) *die $\mathbf{b}_1, \dots, \mathbf{b}_n$ sind linear unabhängig,*
- (ii) *$V = \mathcal{L}(\mathcal{B})$, dh. V ist die lineare Hülle von $\mathbf{b}_1, \dots, \mathbf{b}_n$. Die Teilmenge $\mathbf{b}_1, \dots, \mathbf{b}_r$ mit $r < n$ bildet eine Teilbasis von $\mathcal{L}$.*
- (iii) *Es sei $\mathbf{v} \in V$ und es gelte*

$$\mathbf{v} = a_1 \mathbf{b}_1 + a_2 \mathbf{b}_2 + \dots + a_n \mathbf{b}_n. \quad (2.87)$$

Die Koeffizienten $a_1, \dots, a_n$ heißen Koordinaten von $\mathbf{v}$ bezüglich $\mathcal{B}$ (vergl. Beispiel 2.1, Seite 18).

$V = \mathcal{L}(\mathcal{B})$ bedeutet, dass jeder Vektor aus V als Linearkombination der Basisvektoren $\mathbf{b}_1, \dots, \mathbf{b}_n$ darstellbar ist, d.h. für eine gegebene Basis $\mathcal{B} = \{\mathbf{b}_1, \dots, \mathbf{b}_n\}$ existieren für jeden Vektor $\mathbf{v} \in \mathcal{L}$ Koeffizienten $a_1, \dots, a_n$ (also ein Vektor $\mathbf{a} = (a_1, \dots, a_n)'$) derart, dass die Darstellung eines Vektors $\mathbf{v} \in V$ wie in (2.87) möglich ist. Dies heißt, dass $\mathcal{B}$ eine maximal linear unabhängige Teilmenge von V im Sinne der Definition 2.14 ist. $\mathcal{L}_r = \mathcal{L}(B_r) = \mathcal{L}(\mathbf{b}_1, \dots, \mathbf{b}_r)$ mit $r < n$ definiert einen Teilraum von V (vergl. Satz 2.6).

Definition 2.16 *Es sei V ein Vektorraum mit der Basis $\mathcal{B} = \{\mathbf{b}_1, \dots, \mathbf{b}_n\}$. Dann heißt V n -dimensionaler Vektorraum; man schreibt auch V_n , um die Anzahl der Vektoren in einer Basis von V anzuzeigen. n heißt Dimension des Vektorraums.*

Anmerkung: In Definition 2.16 wird der Begriff des n -dimensionalen Vektorraums durch die Anzahl der Basisvektoren definiert. In der Tat existiert ein Zusammenhang zwischen der Anzahl $n < \infty$ der Komponenten der Vektoren eines Vektorraums V und der Anzahl der Basisvektoren, die notwendig sind, um alle Vektoren von V zu erzeugen. Dieser Zusammenhang wird im Folgenden elaboriert. Zuvor wird aber der Begriff der orthogonalen Basis eingeführt. $\square$

Linear unabhängige Vektoren sind nicht notwendig auch paarweise orthogonal zueinander, aber paarweise orthogonale Vektoren sind notwendig linear unabhängig. Orthogonale Vektoren können demnach als Basisvektoren gewählt werden. Dieser Fall ist besonders wichtig, weshalb eine eigene Definition dafür eingeführt wird:

Definition 2.17 *Es sei V ein n -dimensionaler Vektorraum. Eine Basis*

$$\mathcal{B} = (\mathbf{b}_1, \dots, \mathbf{b}_n)$$

von V heißt Orthonormalbasis (ONB) (oder orthonormale Basis), wenn die $\mathbf{b}_j$ auf die Länge 1 normiert und paarweise orthogonal sind, d.h. wenn

$$\mathbf{b}'_j \mathbf{b}_k = \begin{cases} 0, & j \neq k \\ 1, & j = k \end{cases}, \quad j, k = 1, \dots, n \quad (2.88)$$

Die Basis $\mathcal{B}_r = (\mathbf{b}_1, \dots, \mathbf{b}_r)$ mit $r < n$ heißt orthonormale Teilbasis.

Orthonormale Basisentwicklung eines Vektors: Die zur Darstellung eines beliebigen Vektors $\mathbf{v} \in \mathcal{L}$ benötigten Koeffizienten a_j ergeben sich besonders einfach, wenn Orthonormalbasen gewählt werden: Es sei $\mathbf{x} \in V_n$ ($\mathbf{x}$ sei ein n -dimensionaler Vektor) und die $\mathbf{b}_1, \dots, \mathbf{b}_n$ seien orthonormale Basisvektoren. Dann existieren Koordinaten $a_1, \dots, a_n$ derart, dass

$$\mathbf{x} = a_1 \mathbf{b}_1 + \dots + a_n \mathbf{b}_n = \sum_{k=1}^n a_k \mathbf{b}_k. \quad (2.89)$$

Für die Koeffizienten a_j ergibt sich eine einfache Darstellung. Man betrachte dazu das Skalarprodukt $\mathbf{x}'\mathbf{b}_j$:

$$a_j = \mathbf{x}'\mathbf{b}_j = \sum_{k=1}^n a_k \mathbf{b}'_k \mathbf{b}_j, \quad j = 1, \dots, n \quad (2.90)$$

denn

$$\mathbf{b}'_k \mathbf{b}_j = \begin{cases} 0, & j \neq k \\ 1, & j = k \end{cases} \quad (2.91)$$

(2.89) kann dann in der Form

$$\mathbf{x} = \sum_{k=1}^n (\mathbf{x}'\mathbf{b}_k) \mathbf{b}_k. \quad (2.92)$$

Dieser Ausdruck heißt auch *orthonormale Basisentwicklung* des Vektors $\mathbf{x}$.

Anmerkung: Bekanntlich kann unter bestimmten Normierungsbedingungen ein Skalarprodukt als Korrelation interpretiert werden, dann bedeutet (2.92) $\mathbf{x}'\mathbf{b}_k$ die Korrelation zwischen dem Vektor $\mathbf{x}$ und dem Basisvektor $\mathbf{b}_k$. In der Faktorenanalyse wird eine Ladung einer Variablen auf einer latenten Dimension als Korrelation zwischen dem Item und der latenten Dimension interpretiert. Diese Interpretation beruht auf (2.92). $\square$

Die n -dimensionalen Einheitsvektoren sind ein Beispiel für eine orthonormale Basis:

Definition 2.18 Die n -dimensionalen Einheitsvektoren $\mathbf{e}_1, \dots, \mathbf{e}_n$ mit

$$\mathbf{e}_i = (0, \dots, 0, 1, 0, \dots, 0)'$$

der i -te n -dimensionale Einheitsvektor, bilden eine orthonormale Basis des V_n ; sie heißt die kanonische Basis des V_n .

Die Einheitsvektoren sind linear unabhängig, denn $\vec{0} = \lambda_1 \mathbf{e}_1 + \lambda_2 \mathbf{e}_2 + \cdots + \lambda_n \mathbf{e}_n$ ist nur möglich für $\lambda_1 = \cdots = \lambda_n = 0$; für die i -te Komponente hat man nämlich $0 = \lambda_i 1$, woraus sofort $\lambda_i = 0$ folgt. Darüber hinaus sind die $\mathbf{e}_i$ orthonormal, denn

$$\mathbf{e}_i' \mathbf{e}_j = \begin{cases} 1, & i = j \\ 0, & i \neq j \end{cases}$$

Die Vektoren $\mathbf{e}_1, \dots, \mathbf{e}_n$ bilden deshalb eine orthonormale Basis des V_n .

Man sieht sofort, dass sich jeder Vektor $\mathbf{x} \in V_n$ als Linearkombination der $\mathbf{e}_1, \dots, \mathbf{e}_n$ darstellen läßt (vergl. Beispiel 2.10):

$$\mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} = x_1 \begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix} + x_2 \begin{pmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{pmatrix} + \cdots + x_n \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{pmatrix}. \quad (2.93)$$

Die Komponenten $x_1, \dots, x_n$ von $\mathbf{x}$ heißen auch *Koordinaten von $\mathbf{x}$* bezüglich der Basis $\mathbf{e}_1, \dots, \mathbf{e}_n$.

Satz 2.9 *Es sei $V = V_n$ ein n -dimensionaler Vektorraum und $\mathcal{B}_r = \{\mathbf{b}_1, \dots, \mathbf{b}_r\}$, $r < n$ eine Teilbasis von V . Dann ist die lineare Hülle $U = \mathcal{L}(\mathcal{B}_r)$ von $\mathcal{B}_r$ ein Teilvektorraum von V .*

Beweis: Der Satz kennzeichnet einen Spezialfall von Satz 2.6, Seite 45, aber ein gesonderter Beweis kann nicht schaden. Es seien $\mathbf{v} = a_1 \mathbf{b}_1 + \cdots + a_r \mathbf{b}_r$ und $\mathbf{u} = b_1 \mathbf{b}_1 + \cdots + b_r \mathbf{b}_r$ Linearkombinationen der Vektoren aus der Teilbasis; dann sind $\mathbf{v}, \mathbf{u} \in U$. Dann ist

$$\mathbf{u} + \mathbf{v} = (a_1 + b_1) \mathbf{b}_1 + \cdots + (a_r + b_r) \mathbf{b}_r \in U,$$

denn $\mathbf{v} + \mathbf{u}$ ist offenbar ebenfalls eine Linearkombination der Vektoren aus $\mathcal{B}$. Ebenso ist $\lambda \mathbf{u} \in U$, da $\lambda \mathbf{u}$ wieder eine Linearkombination der $\mathbf{b}_j$, $j = 1, \dots, r$ ist, – dies gilt für alle $\mathbf{u} \in U$. $\square$

Es sei eine Teilmenge $S = \{\mathbf{x}_1, \dots, \mathbf{x}_k\}$ mit $k < n$ Vektoren gegeben. $\mathcal{L}(S)$ ist ein Teilraum von V_n . Sind die $\mathbf{x}_j \in S$ linear abhängig, so existiert eine Teilbasis $\mathcal{B}_r = \{\mathbf{b}_1, \dots, \mathbf{b}_r\}$ mit $r < k$ und $\mathcal{L}(S) = \mathcal{L}(\mathcal{B}_r)$. Die folgende Definition führt im Wesentlichen eine Redeweise ein:

Definition 2.19 *Es sei V ein Vektorraum und S eine Teilmenge von V . Dann ist der Rang von S gleich der Dimension des von S erzeugten Unterraums $\mathcal{L}(S)$. Ist $V = V_n$ ein n -dimensionaler Vektorraum und ist $S \subset V_n$, so hat S den Rang $r < n$, wenn S r linear unabhängige Vektoren enthält; für $r = n$ hat S den vollen Rang.*

Anmerkung: r heißt auch die *Dimension* des Unterraums $\mathcal{L}(S)$, und $n - r$ heißt die *Kodimension* des Unterraums $\mathcal{L}(S)$. Die Dimension eines Unter- oder Teilraums eines Vektorraums ist also nicht notwendig gleich der Dimension, d.h. der Anzahl der Komponenten der Vektoren, die die Elemente des Teilraums sind. $\square$

Die folgenden Betrachtungen werden durch die Frage motiviert, wie viele Vektoren eine Teilmenge M eines Vektorraums haben muß, damit M eine Basis bildet. Dazu wird zunächst der Begriff eines Erzeugendensystems eingeführt:

Definition 2.20 *Es sei V' ein Teilraum des Vektorraums V . V' wird durch die Teilmenge M von Vektoren aus V erzeugt, wenn für die lineare Hülle $\mathcal{L}(M)$ die Aussage $\mathcal{L}(M) = V'$ gilt. M heißt dann lineares Erzeugendensystem von V' . Der Vektorraum V heißt endlich erzeugt, wenn V ein lineares Erzeugendensystem enthält, das nur aus endlich vielen Elementen besteht. M heißt minimales Erzeugendensystem, wenn kein Vektor aus M entfernt werden darf, damit $\mathcal{L}(M) = V'$ gilt (wird also ein Element aus M entfernt, so ist $\mathcal{L}(M) \subset V'$).*

Ist ein Vektorraum endlich erzeugt, so heißt dies, dass es eine endliche Menge $\{\mathbf{v}_1, \dots, \mathbf{v}_n\}$ gibt derart, dass V die lineare Hülle $\mathcal{L}(\{\mathbf{v}_1, \dots, \mathbf{v}_n\})$ von $\{\mathbf{v}_1, \dots, \mathbf{v}_n\}$ ist, d.h.

$$V = \mathcal{L} \left\{ \sum_{j=1}^n a_j \mathbf{v}_j \mid a_j \in \mathbb{K} \right\}, \quad (2.94)$$

wobei hier im Allgemeinen $\mathbb{K} = \mathbb{R}$ oder $\mathbb{K} = \mathbb{C}$. Ist dagegen $V = C^0(\mathbb{R})$ der Vektorraum aller auf $\mathbb{R}$ stetigen Funktionen, so ist V *nicht* endlich erzeugt (hier ohne Beweis).

Für eine Basis eines Vektorraums können die folgenden Aussagen gemacht werden:

Satz 2.10 *M sei eine Teilmenge des Vektorraums V . Die folgenden Aussagen sind äquivalent (im Sinne von $(1) \Rightarrow (2)$ und $(2) \Rightarrow (1)$, etc):*

- (1) *M ist eine Basis von V*
- (2) *M ist ein minimales Erzeugendensystem von V*
- (3) *M ist eine maximal linear unabhängige Teilmenge von V .*

Beweis: $(1) \Rightarrow (2)$ M ist eine Basis $\Rightarrow M$ ist l.u. und $\mathcal{L}(M) = V$. Es sei $A \subset M$ eine echte Teilmenge von M und $M \setminus A$ ($= M$ minus A , das Zeichen $\setminus$ ersetzt das Minuszeichen bei Mengendifferenzen) die Menge von Vektoren aus M , die nicht zu A gehören. Sei weiter $\mathbf{v} \in M \setminus A$. Wegen der linearen Unabhängigkeit der Elemente in V kann $\mathbf{v}$ nicht als Linearkombination der Elemente in $M \setminus A$ dargestellt werden, so dass $\mathcal{L}(M \setminus A) \subset V$ und $\mathcal{L}(A) \subset V$, so dass (2) folgt, M ist ein minimales Erzeugendensystem.

$(2) \Rightarrow (3)$ Es sei M ein minimales Erzeugendensystem von V . Es werde angenommen, dass M linear abhängig ist. Dann existiert ein $\mathbf{v} \in M$ und $\mathbf{v}$ ist

Linearkombination der Elemente von $M \setminus \{\mathbf{v}\}$. Aber dann ist M kein minimales Erzeugendensystem, im Widerspruch zur Annahme, dass M ein minimales Erzeugendensystem ist. Also muß M linear unabhängig sein. Darüber hinaus muß M maximal linear unabhängig sein, denn M ist ein Erzeugendensystem für V , so dass $M \cup \{\mathbf{v}\}$ linear abhängig ist für jedes Element $\mathbf{v} \in V \setminus M$.

(3) $\Rightarrow$ (1) M sei eine maximal linear unabhängige Teilmenge von V . Dann ist für $\mathbf{v} \notin M$, $\mathbf{v} \in V$ die Vereinigung $M \cup \{\mathbf{v}\}$ linear abhängig (denn sonst wäre M nicht *maximal* linear unabhängig). Mithin existieren $\mathbf{v}_1, \dots, \mathbf{v}_n \in M$ und $a, a_1, \dots, a_n \in \mathbb{K}$, die nicht alle gleich Null sind, und es gilt

$$a\mathbf{v} + a_1\mathbf{v}_1 + \dots + a_n\mathbf{v}_n = \vec{0}$$

mit $A \neq 0$, denn sonst wären die $\mathbf{v}_1, \dots, \mathbf{v}_n$ nicht linear unabhängig. Mithin existiert für $\mathbf{v}$ die Darstellung

$$\mathbf{v} = -\frac{a_1}{a}\mathbf{v}_1 - \dots - \frac{a_n}{a}\mathbf{v}_n$$

als Linearkombination der Elemente von M , d.h. M ist eine Basis von V . $\square$

Es läßt sich zeigen, dass jeder Vektorraum V eine Basis hat. Wird V nicht endlich erzeugt, erfordert der Beweis die Anwendung von Resultaten der Mengenlehre, worauf hier nicht eingegangen werden kann. Die meisten multivariaten Techniken beziehen sich aber auf endlich erzeugte Vektorräume, für die der Beweis nicht schwierig ist. Denn 'endlich erzeugt' heißt ja nur, dass eine maximal linear unabhängige Teilmenge M von V existiert, und nach Satz 2.10 folgt daraus die Existenz einer Basis.

Die Dimension eines Vektorraums ist die Anzahl der Basisvektoren, die notwendig sind, um alle Vektoren des Vektorraums als Linearkombination der Basisvektoren zu erzeugen. Es gilt der *Austauschsatz von Steiner*:

Satz 2.11 *Es sei V ein Vektorraum mit der Basis $\mathcal{B} = (\mathbf{v}_1, \dots, \mathbf{v}_n)$. Für $\vec{0} \neq \mathbf{v} \in V$ gibt es dann die Darstellung*

$$\mathbf{v} = a_1\mathbf{v}_1 + \dots + a_n\mathbf{v}_n \tag{2.95}$$

Jeder Vektor $\mathbf{v}_j \in \mathcal{B}$ mit $a_j \neq 0$ kann durch $\mathbf{v}$ ausgetauscht werden, und

$$\{\mathbf{v}_1, \mathbf{v}_{j-1}, \mathbf{v}, \mathbf{v}_{j+1}, \dots, \mathbf{v}_n\} \tag{2.96}$$

ist ebenfalls eine Basis von V .

Beweis: Da $\mathbf{v} \neq \vec{0}$ sind nicht alle $a_j = 0$. Da die Anordnung der $\mathbf{v}_j$ in $\mathcal{B}$ beliebig ist, kann angenommen werden, dass $a_1 \neq 0$, so dass $\mathbf{v}_1$ durch $\mathbf{v}$ ersetzt werden kann. Dann gilt

$$b_1\mathbf{v} + b_2\mathbf{v}_2 + \dots + b_n\mathbf{v}_n = \vec{0}.$$

Da zu zeigen ist, dass $\{\mathbf{v}, \mathbf{v}_2, \dots, \mathbf{v}_n\}$ eine Basis ist, muß gezeigt werden, dass die $\mathbf{v}, \mathbf{v}_2, \dots, \mathbf{v}_n$ linear unabhängig sind, d.h. die Darstellung des Vektors $\vec{0}$ nur möglich ist, wenn $b_1 = \dots = b_n = 0$. Setzt man den obigen Ausdruck der Linearkombination von $\mathbf{v}$ ein und vereinfacht, so erhält man den Ausdruck

$$b_1 a_1 \mathbf{v}_1 + (b_1 a_2 + b_2) \mathbf{v}_2 + \dots + (b_1 a_n + b_n) \mathbf{v}_n = \vec{0}.$$

Wegen der linearen Unabhängigkeit der $\mathbf{v}_j$ folgt dann aber

$$b_1 a_1 = b_1 a_2 + b_2 = \dots = b_1 a_n + b_n = 0.$$

Da $a_1 \neq 0$ folgt $b_1 = 0$ und also $b_1 = \dots = b_n = 0$, d.h. die Vektoren $\{\mathbf{v}, \mathbf{v}_2, \dots, \mathbf{v}_n\}$ sind linear unabhängig.

Jetzt muß nur noch gezeigt werden, dass $\{\mathbf{v}, \mathbf{v}_2, \dots, \mathbf{v}_n\}$ ein Erzeugendensystem für V ist. Dazu sei $\mathbf{w} \in V$ ein beliebiger Vektor aus V . Da $\{\mathbf{v}_1, \dots, \mathbf{v}_n\}$ nach Voraussetzung eine Basis ist, existieren Koeffizienten $c_1, \dots, c_n \in \mathbb{K}$ derart, dass

$$\mathbf{w} = c_1 \mathbf{v}_1 + \dots + c_n \mathbf{v}_n.$$

Aus (2.95) folgt

$$\mathbf{v}_1 = \frac{1}{a_1} (\mathbf{v} - a_2 \mathbf{v}_2 - \dots - a_n \mathbf{v}_n)$$

Setzt man diesen Ausdruck für $\mathbf{v}_1$ in den Ausdruck für $\mathbf{w}$ ein, so sieht man, dass $\mathbf{w}$ als Linearkombination der $\{\mathbf{v}, \mathbf{v}_2, \dots, \mathbf{v}_n\}$ dargestellt werden kann, so dass diese Vektoren ebenfalls eine Basis für V bilden. $\square$

Satz 2.12 *Die Darstellung $V = \mathcal{L}(\mathcal{B})$ ist für eine gegebene Basis $\mathcal{B}$ eindeutig. Die Wahl einer Basis $\mathcal{B}$ für einen Vektorraum V ist nicht eindeutig.*

Beweis: Angenommen, es gäbe zwei Darstellungen der Form

$$\mathbf{v} = a_1 \mathbf{b}_1 + \dots + a_n \mathbf{b}_n = b_1 \mathbf{b}_1 + \dots + b_n \mathbf{b}_n.$$

Dann folgt

$$(a_1 - b_1) \mathbf{b}_1 + \dots + (a_n - b_n) \mathbf{b}_n = 0$$

und wegen der linearen Unabhängigkeit der Basisvektoren $\mathbf{b}_j$ muß $a_j - b_j = 0$ für $j = 1, \dots, n$ gelten, d.h. aber $a_j = b_j$ für alle j , V wird eindeutig durch $\mathcal{B}$ bestimmt.

Die Wahl einer Basis für einen Vektorraum ist *nicht* eindeutig. So sei $\mathcal{B}_1 = \{\mathbf{b}_1, \dots, \mathbf{b}_n\}$ eine Basis des V_n , und $\mathbf{c}_1$ sei eine Linearkombination der Elemente von $\mathcal{B}_1$. Nach dem Austauschsatz kann etwa $\mathbf{b}_1$ durch $\mathbf{c}_1$ ersetzt werden, so dass die Basis $\mathcal{B}_2 = \{\mathbf{c}_1, \mathbf{b}_2, \dots, \mathbf{b}_n\}$ entsteht. $\mathbf{c}_2$ sei eine Linearkombination der $\mathbf{c}_1, \mathbf{b}_2, \dots, \mathbf{b}_n$. Hier kann $\mathbf{b}_2$ durch $\mathbf{c}_2$ ausgetauscht werden, etc. Auf diese Weise entsteht die Basis $\mathcal{C} = \{\mathbf{c}_1, \dots, \mathbf{c}_n\}$. Andererseits sind die $\mathbf{c}_j$, $j = 1, \dots, n$ Elemente der linearen Hülle $\mathcal{L}(\mathcal{B}_1)$ von $\mathcal{B}_1$, und da die $\mathbf{c}_j$ eine Basis bilden, sind sie linear

unabhängig. Damit ist gezeigt, dass die lineare Hülle einer Basis $\mathcal{B}$ mindestens eine Menge $\mathbf{c}_1, \dots, \mathbf{c}_n$ enthält, die ebenfalls eine Basis $\mathcal{C}$ bilden, die ebenso gut wie $\mathcal{B}$ als Basis gewählt werden können. Da $\mathbf{c}_1 = a_1 \mathbf{b}_1 + \dots + a_n \mathbf{b}_n$ und die Koeffizienten $a_1, \dots, a_n$ beliebig gewählt werden können, ebenso die Koeffizienten für $\mathbf{c}_2$ etc kann man folgern, dass $\mathcal{L}(\mathcal{B})$ beliebig viele Basen enthält, die ebenso wie $\mathcal{B}$ gewählt werden können. $\square$

Beispiel 2.17 Es sei $V = \mathbb{R}^2$ der Vektorraum aller 2-dimensionalen Vektoren. Dann ist die Teilmenge

$$\left\{ \begin{pmatrix} a \\ b \end{pmatrix}, \begin{pmatrix} c \\ d \end{pmatrix} \right\}$$

eine Basis für V , falls die Bedingung $ad - bc \neq 0$ erfüllt ist; dazu müssen die Vektoren $(a, b)'$ und $(c, d)'$ nicht orthogonal sein. Angenommen, es sei $ad - bc = 0$, so dass $ad = -bc$, woraus $a/b = -c/d = q$ folgt, was $a = qb$ und $c = -qc$ impliziert, d.h. die beiden Vektoren $(a, b)'$ und $(c, d)'$ liegen auf der gleichen Geraden und können deshalb nicht *alle* 2-dimensionalen Vektoren erzeugen. $\square$

Man kann nun fragen, welche Beziehung zwischen den Koeffizienten u_j und v_j besteht. Man macht sich leicht klar, dass es einen beträchtlichen Rechenaufwand bedeutet, diese Beziehung "elementar" herzuleiten, – weshalb dieser Ansatz hier auch nicht weiter verfolgt wird. Für Vektoren $\mathbf{v} \in \mathbb{R}^n$ kann mit den Mitteln des Matrixkalküls sehr schnell eine Antwort auf die Frage nach dieser Beziehung gegeben werden, vergl. Beispiel 3.7, Seite 94.

Satz 2.13 *Der Vektorraum V habe eine Basis $\mathcal{B} = (\mathbf{v}_1, \dots, \mathbf{v}_n)$. Dann ist jede Teilmenge M mit $m > n$ Elementen linear abhängig. Zwei verschiedene Basen von V haben stets dieselbe Anzahl von Elementen aus V .*

Anmerkung: Dieser Satz wird gelegentlich als *Fundamental-Lemma der Linearen Algebra* bezeichnet (etwa in Koecher (1997), p. 21). $\square$

Beweis: Es sei $\mathcal{B} = (\mathbf{v}_1, \dots, \mathbf{v}_n)$ eine Basis von V , und es sei $\{\mathbf{w}_1, \dots, \mathbf{w}_m\} \subset V$, $m > n$. Es werde angenommen, die $\mathbf{w}_1, \dots, \mathbf{w}_m$ seien linear unabhängig. Nach dem Austauschsatz 2.11 kann einer der Basisvektoren aus $\mathcal{B}$ etwa durch den Vektor $\mathbf{v} = \mathbf{w}_1$ ausgetauscht werden. Man erhält dann die Basis $\mathbf{w}_1, \mathbf{v}_2, \dots, \mathbf{v}_n$, und $\mathbf{w}_2$ kann als Linearkombination

$$\mathbf{w}_2 = b_1 \mathbf{w}_1 + a_2 \mathbf{v}_2 + \dots + a_n \mathbf{v}_n$$

ausgedrückt werden. So fährt man weiter fort, indem man $\mathbf{v}_2$ durch $\mathbf{w}_2$ austauscht, etc, so dass man schließlich die $\mathbf{v}_1, \dots, \mathbf{v}_n$ durch die $\mathbf{w}_1, \dots, \mathbf{w}_n$ ersetzt hat. Aber $m > n$, und $\mathbf{w}_m$ kann nun als Linearkombination der $\mathbf{w}_1, \dots, \mathbf{w}_n$ ausgedrückt werden. Aber das heißt, dass die $\mathbf{w}_1, \dots, \mathbf{w}_m$ insgesamt linear abhängig sind, entgegen der Annahme ihrer linearen Unabhängigkeit. Es können nur n der m

Elemente $\mathbf{w}_j$ linear unabhängig sein, und das heißt, dass alle Basen von V dieselbe Anzahl von Elementen haben müssen. $\square$

Schreibweisen: Für die Dimension von V wird $\dim_{\mathbb{K}} V$ oder einfach $\dim V$ geschrieben, wenn klar ist, um welchen Körper $\mathbb{K}$ es sich handelt. Für einen endlich erzeugten Vektorraum gilt $\dim_{\mathbb{K}} V < \infty$, für einen nicht endlich erzeugten Vektorraum gilt $\dim_{\mathbb{K}} V = \infty$.

Aus dem Vorangegangenen folgt für einen n -dimensionalen Vektorraum V_n :

1. Ein Erzeugendensystem für V_n besteht aus n Vektoren.
2. Mehr als n Vektoren des V_n sind stets linear abhängig.
3. Ein Erzeugendensystem mit n Vektoren bildet eine Basis für V_n .

Übung: In Satz 2.1, Seite 32 wurde gezeigt, dass für p linear unabhängige n -dimensionale Vektoren $\mathbf{x}_1, \dots, \mathbf{x}_p$ die Darstellung des n -dimensionalen Vektors $\mathbf{y} = \lambda_1 \mathbf{x}_1 + \dots + \lambda_p \mathbf{x}_p$ eindeutig ist, d.h. es gibt nur einen Vektor $\vec{\lambda} = (\lambda_1, \dots, \lambda_p)'$, der dieser Beziehung genügt. Die Frage war, ob $\vec{\lambda}$ nicht mehr eindeutig ist, wenn die $\mathbf{x}_1, \dots, \mathbf{x}_p$ linear abhängig sind, d.h. ob die lineare Unabhängigkeit der $\mathbf{x}_j$ nicht nur hinreichend, sondern auch notwendig für die Eindeutigkeit von $\vec{\lambda}$ ist. Es werde nun angenommen, dass die $\mathbf{x}_j$ linear abhängig sind. Dann existieren Basisvektoren $\mathbf{v}_1, \dots, \mathbf{v}_r$, $r < p$, und Koeffizienten a_{kj} , $k = 1, \dots, r$, mit

$$\mathbf{x}_j = a_{1j} \mathbf{v}_1 + \dots + a_{rj} \mathbf{v}_r,$$

so dass

$$\mathbf{y} = \lambda_1 (a_{11} \mathbf{v}_1 + \dots + a_{r1} \mathbf{v}_r) + \dots + \lambda_p (a_{p1} \mathbf{v}_1 + \dots + a_{pr} \mathbf{v}_r)$$

gilt; umgeordnet ergibt sich

$$\mathbf{y} = \underbrace{(\lambda_1 a_{11} + \dots + \lambda_p a_{p1})}_{b_1} \mathbf{v}_1 + \dots + \underbrace{(\lambda_1 a_{1r} + \dots + \lambda_p a_{pr})}_{b_r} \mathbf{v}_r.$$

Da die $\mathbf{v}_j$ nach Voraussetzung linear unabhängig sind, folgt aus Satz 2.1, dass die $b_1, \dots, b_r$ eindeutig sind. Wenn nun $\vec{\lambda}$ nicht eindeutig ist, so existiert ein Vektor $\vec{\mu} = (\mu_1, \dots, \mu_p)'$ derart, dass auch

$$\mathbf{y} = \underbrace{(\mu_1 a_{11} + \dots + \mu_p a_{p1})}_{b_1} \mathbf{v}_1 + \dots + \underbrace{(\mu_1 a_{1r} + \dots + \mu_p a_{pr})}_{b_r} \mathbf{v}_r$$

gilt. Subtrahiert man die beiden Ausdrücke für $\mathbf{y}$, so erhält man eine Darstellung des Nullvektors

$$\begin{aligned} \vec{0} &= ((\lambda_1 - \mu_1) a_{11} + \dots + (\lambda_p - \mu_p) a_{p1}) \mathbf{v}_1 + \\ &\quad \dots + ((\lambda_1 - \mu_1) a_{1r} + \dots + (\lambda_p - \mu_p) a_{pr}) \mathbf{v}_r. \end{aligned}$$

Wegen der linearen Unabhängigkeit der $\mathbf{v}_j$ folgt nun aber, dass die Koeffizienten der $\mathbf{v}_j$ gleich Null sein müssen. Die Vektorgleichung entspricht also dem folgenden System von Gleichungen:

$$\begin{aligned}(\lambda_1 - \mu_1)a_{11} + (\lambda_2 - \mu_2)a_{12} + \cdots + (\lambda_p - \mu_p)a_{p1} &= 0 \\(\lambda_1 - \mu_1)a_{12} + (\lambda_2 - \mu_2)a_{12} + \cdots + (\lambda_p - \mu_p)a_{p2} &= 0 \\&\vdots \\(\lambda_1 - \mu_1)a_{1r} + (\lambda_2 - \mu_2)a_{12} + \cdots + (\lambda_p - \mu_p)a_{pr} &= 0\end{aligned}$$

das in der Form

$$(\lambda_1 - \mu_1)\mathbf{a}_1 + \cdots + (\lambda_p - \mu_p)\mathbf{a}_p = \vec{0}$$

geschrieben werden kann. Darin sind die $\mathbf{a}_j = (a_{j1}, \dots, a_{jr})'$, $j = 1, \dots, p$ r -dimensionale Vektoren. Da $p > r$ folgt nach Satz 2.13, dass die $\mathbf{a}_j$ linear abhängig sind, so dass die $(\lambda_j - \mu_j)$ nicht alle gleich Null sind. Dies bedeutet aber, dass $\boldsymbol{\lambda}$ nicht eindeutig ist, d.h. es gibt verschiedene Vektoren $\boldsymbol{\lambda}, \boldsymbol{\mu}, \dots$, die der Darstellung von $\mathbf{y}$ durch die $\mathbf{x}_1, \dots, \mathbf{x}_p$ genügen. Dies wiederum bedeutet, dass die lineare Unabhängigkeit der $\mathbf{x}_j$ auch notwendig für die Eindeutigkeit der Darstellung von $\mathbf{y}$ ist. $\square$

Ein Begriff, der sich gelegentlich als nützlich erweist, ist der des orthogonalen Komplements einer Teilmenge von Vektoren eines Vektorraums:

Definition 2.21 *Es sei V ein Vektorraum und $M \subseteq V$ sei eine Teilmenge von Vektoren aus V . Dann heißt¹³*

$$M^\perp := \{\mathbf{w} \in V \mid \mathbf{w}'\mathbf{v} = 0 \text{ für alle } \mathbf{v} \in M\} \quad (2.97)$$

das orthogonale Komplement von M in V .

Die Menge M muß kein Teilraum von V sein. Es zeigt sich aber, daß $M^\perp$ stets ein Teilraum von V ist:

Satz 2.14 *Es sei V ein Vektorraum. Dann gelten die folgenden Aussagen:*

- (a) *Es seien M_1 und M_2 Teilmengen von V mit $M_1 \subseteq M_2$. Dann gilt $M_2^\perp \subseteq M_1^\perp$.*
- (b) *Sei M irgendeine Teilmenge von V ; dann ist $M^\perp$ ein Teilraum von V .*
- (c) *Sei M irgendeine Teilmenge von V ; dann ist $M^\perp = \mathcal{L}(M^\perp)$.*
- (d) *Es seien $\mathbf{v}_1, \dots, \mathbf{v}_k$ beliebige Vektoren aus V . Dann gilt*

$$\mathcal{L}((\mathbf{v}_1, \dots, \mathbf{v}_k)^\perp) = \mathcal{L}((\mathbf{v}_1)^\perp \cap \cdots \cap \mathcal{L}((\mathbf{v}_k)^\perp)).$$

¹³Das Zeichen $\perp$ steht allgemein für orthogonal, senkrecht auf, etc.

Beweis: (a) folgt sofort aus der Definition 2.21. (b) Es seien $\mathbf{u}, \mathbf{v} \in M^\perp$ und $\lambda, \mu \in \mathbb{R}$ (allgemein: $\lambda, \mu \in \mathbb{K}$). Dann folgt sofort $\lambda\mathbf{u} + \mu\mathbf{v} \in M^\perp$, da offenbar $(\lambda\mathbf{u} + \mu\mathbf{v})'\mathbf{w} = 0$ für jedes Element $\mathbf{w} \in M$. (c) Aus (a) folgt $\mathcal{L}(M^\perp) \subseteq M^\perp$. Andererseits gibt es für jeden Vektor $\mathbf{w} \in M$ Vektoren $\mathbf{w}_1, \dots, \mathbf{w}_k$ und Skalare $\lambda_1, \dots, \lambda_k$ derart, dass $\mathbf{w} = \lambda_1\mathbf{w}_1 + \dots + \lambda_k\mathbf{w}_k$, und für beliebigen Vektor $\mathbf{v} \in M^\perp$ gilt $\mathbf{v}'\mathbf{w}$, da notwendig $\mathbf{v}'\mathbf{w}_1 = \dots = \mathbf{v}'\mathbf{w}_k = 0$ gilt. Dann gilt auch für jede Linearkombination $\mathbf{v} = a\mathbf{v}_1 + b\mathbf{v}_2$ von Vektoren $\mathbf{v}_1$ und $\mathbf{v}_2$ aus $M^\perp$, dass $\mathbf{v}'\mathbf{w} = 0$. (d) Es sei $M = \{\mathbf{v}_1, \dots, \mathbf{v}_k\}$ und $\mathbf{v} \in \mathcal{L}(M)$; dann muß $\mathbf{v}'\mathbf{v}_j = 0$ gelten für alle $j = 1, \dots, k$. Dies heißt aber, dass $\mathbf{v} \in \mathcal{L}(\mathbf{v}_1) \cap \dots \cap \mathcal{L}(\mathbf{v}_k)$. Es seien $\lambda_1, \dots, \lambda_k \in \mathbb{K}$ ($\mathbb{R}$), und $\mathbf{w} = \lambda_1\mathbf{w}_1 + \dots + \lambda_k\mathbf{w}_k$. Dann muß auch $\mathbf{v}'\mathbf{w} = 0$ gelten, da ja $\mathbf{v}'\mathbf{w}_j = 0$ für alle j , woraus $\mathbf{v} \in M^\perp$ folgt. $\square$

Der folgende Satz erweist sich für viele Betrachtungen als nützlich, insbesondere wenn Projektionen von Vektoren betrachtet werden (vergl. Abschnitt 3.17).

Satz 2.15 *Es seien V und W Teilräume des $\mathbb{R}^n$ mit der Eigenschaft $W = V^\perp$. Dann existieren für einen beliebigen $\mathbf{x} \in \mathbb{R}^n$ Vektoren $\mathbf{v} \in V$ und $\mathbf{w} \in W$ derart, dass*

$$\mathbf{x} = \mathbf{v} + \mathbf{w}, \quad (2.98)$$

wobei die Vektoren $\mathbf{v}$ und $\mathbf{w}$ eindeutig bestimmt sind.

Beweis: Es sei $(\mathbf{a}_1, \dots, \mathbf{a}_k)$ eine orthonormale Basis von V . Es sei $\mathbf{v} = \sum_{i=1}^k c_i \mathbf{a}_i$; dann ist $\mathbf{v} \in V$, da $\mathbf{v}$ ja als Linearkombination der Basisvektoren von V definiert ist. Insbesondere seien die c_i als Skalarprodukte $c_i = \mathbf{x}'\mathbf{a}_i$ definiert. Weiter sei $\mathbf{w} = \mathbf{x} - \mathbf{v}$. Dann folgt

$$\mathbf{w}'\mathbf{a}_k = \mathbf{x}'\mathbf{a}_k - \mathbf{v}'\sum_{i=1}^k c_i \mathbf{a}_i' \mathbf{a}_k = c_k - c_k = 0,$$

da $\mathbf{a}_i'\mathbf{a}_k = 0$ für $i \neq k$ und $\mathbf{a}_i'\mathbf{a}_k = 1$ für $i = k$. Also muß $\mathbf{w} \in V^\perp = W$ sein. Damit ist gezeigt, dass eine Darstellung von $\mathbf{x}$ der Form (2.98) existiert. Nun muß noch gezeigt werden, dass diese Darstellung eindeutig ist. Dazu werde angenommen, dass eine zweite Darstellung $\mathbf{x} = \mathbf{v}^* + \mathbf{w}^*$ existiert. Dann folgt

$$\mathbf{x} - \mathbf{x} = \mathbf{v} + \mathbf{w} - \mathbf{v}^* - \mathbf{w}^* = 0.$$

Es sei $\tilde{\mathbf{v}} = \mathbf{v} - \mathbf{v}^* \in V$, $\tilde{\mathbf{w}} = \mathbf{w} - \mathbf{w}^* \in W$, und es gilt $\tilde{\mathbf{v}} + \tilde{\mathbf{w}} = 0$, so dass $\tilde{\mathbf{v}} = -\tilde{\mathbf{w}}$. Dann muß aber $\tilde{\mathbf{v}}'\tilde{\mathbf{w}} = -\tilde{\mathbf{w}}'\tilde{\mathbf{w}} = 0$ gelten, wegen $\tilde{\mathbf{v}} \in V$ und $\tilde{\mathbf{w}} \in V^\perp = W$. $\tilde{\mathbf{w}}'\tilde{\mathbf{w}} = 0$ ist aber nur möglich, wenn $\tilde{\mathbf{w}} = \mathbf{w} - \mathbf{w}^* = \vec{0}$, d.h. $\mathbf{w}$ ist eindeutig bestimmt. Analog folgt $\tilde{\mathbf{v}}'\tilde{\mathbf{v}} = -\tilde{\mathbf{w}}'\tilde{\mathbf{v}} = 0$, so dass $\tilde{\mathbf{v}} = \vec{0}$ und damit die Eindeutigkeit von $\mathbf{v}$ folgt. $\square$

Elementare Umformungen: Es seien $\mathbf{v}_1, \dots, \mathbf{v}_n$ Vektoren; die Maximalzahl linear unabhängiger Vektoren unter ihnen ist der *Rang* dieser Menge von Vektoren. Die folgenden Operationen auf der Vektorenmenge verändern den Rang nicht:

- (1) Vertauschung zweier Vektoren:
 $\cdot \quad (\mathbf{v}_1, \dots, \mathbf{v}_i, \dots, \mathbf{v}_j, \dots, \mathbf{v}_n) \mapsto (\mathbf{v}_1, \dots, \mathbf{v}_j, \dots, \mathbf{v}_i, \dots, \mathbf{v}_n)$
- (2) Multiplikation eines Vektors mit einem Skalar λ :
 $\cdot \quad (\mathbf{v}_1, \dots, \mathbf{v}_i, \dots, \mathbf{v}_n) \mapsto (\mathbf{v}_1, \dots, \lambda \mathbf{v}_i, \dots, \mathbf{v}_n)$
- (3) Ersetzung eines Vektors $\mathbf{v}_i$ durch $\mathbf{v}_i + \lambda \mathbf{v}_j$, $i \neq j$, $\lambda \in \mathbb{R}$.

Entsteht durch elementare Umordnungen aus $S = \{\mathbf{v}_1, \dots, \mathbf{v}_n\}$ die Menge $S' = \{\mathbf{v}'_1, \dots, \mathbf{v}'_n\}$, so ist $\mathcal{L}(S) = \mathcal{L}(S')$ und S und S' haben denselben Rang.

Satz 2.16 (*Dimensionssatz*) *Es sei W ein Vektorraum mit der Dimension $n = \dim(W) < \infty$, und $U \subseteq W$, $V \subseteq W$ seien Teilräume von W . Dann gilt*

$$\dim(U) + \dim(V) = \dim(U \cap V) + \dim(U + V), \quad (2.99)$$

bzw.

$$\dim(U + V) = \dim(U) + \dim(V) - \dim(U \cap V). \quad (2.100)$$

Beweis: Es sei $D = U \cap V$. Da U und V Teilräume sind, ist D ebenfalls ein Teilraum (Satz 2.4). D habe die Basis $B_D = (\mathbf{d}_1, \dots, \mathbf{d}_r)$. Die $\mathbf{d}_j$ sind sowohl aus U als auch aus V . Durch Hinzunahme geeigneter Vektoren aus U bzw. V läßt sich B_D zu einer Basis für U erweitern: $B_U = (\mathbf{d}_1, \dots, \mathbf{d}_r, \mathbf{u}_1, \dots, \mathbf{u}_s)$, und ebenso läßt sich B_D zu einer Basis für V erweitern: $B_V = (\mathbf{d}_1, \dots, \mathbf{d}_r, \mathbf{v}_1, \dots, \mathbf{v}_t)$. Es ist also $\dim(D) = \dim(U \cap V) = r$, $\dim(U) = r + s$ und $\dim(V) = r + t$. Weiter sei $S = U + V$. Zu zeigen ist nun die Behauptung (2.99), d.h. dass die Basis von S durch

$$\dim(S) = B_S = ((\mathbf{d}_1, \dots, \mathbf{d}_r, \mathbf{u}_1, \dots, \mathbf{u}_s, \mathbf{v}_1, \dots, \mathbf{v}_t))$$

gegeben ist.

Es sei $\mathbf{s} \in S$. Dann existieren Vektoren $\mathbf{u} \in U$, $\mathbf{v} \in V$ derart, dass $\mathbf{s} = \mathbf{u} + \mathbf{v}$. $\mathbf{u}$ und $\mathbf{v}$ lassen sich als Linearkombinationen der jeweiligen Basisvektoren darstellen:

$$\begin{aligned} \mathbf{u} &= \sum_{i=1}^r \alpha_i \mathbf{d}_i + \sum_{i=1}^s \lambda_i \mathbf{u}_i \\ \mathbf{v} &= \sum_{i=1}^r \beta_i \mathbf{d}_i + \sum_{j=1}^t \mu_j \mathbf{v}_j \end{aligned}$$

$\alpha_i, \beta_i, \lambda_i, \mu_i \in \mathbb{K}$ (hier $\mathbb{K} = \mathbb{R}$). Daraus folgt

$$\mathbf{s} = \sum_{i=1}^r \alpha_i \mathbf{d}_i + \sum_{i=1}^s \lambda_i \mathbf{u}_i + \sum_{i=1}^r \beta_i \mathbf{d}_i + \sum_{i=1}^t \mu_i \mathbf{v}_i,$$

und $\mathbf{s}$ ist eine Linearkombination der Vektoren aus B_S .

Jetzt muß noch gezeigt werden, dass die Vektoren aus B_S linear unabhängig sind. Dazu wird $\vec{0}$ als Linearkombination dieser Vektoren dargestellt, und im Fall linearer Unabhängigkeit müssen die Koeffizienten gleich Null sein:

$$\vec{0} = \sum_{i=1}^r \delta_i \mathbf{d}_i + \sum_{i=1}^s \lambda_i \mathbf{u}_i + \sum_{i=1}^t \mu_i \mathbf{w}_i,$$

woraus

$$-\sum_{i=1}^t \mu_i \mathbf{w}_i = \sum_{i=1}^r \delta_i \mathbf{d}_i + \sum_{i=1}^s \lambda_i \mathbf{u}_i$$

folgt. Aber es ist sicherlich

$$\sum_{i=1}^t \mu_i \mathbf{w}_i \in D = U \cap V,$$

so dass es Koeffizienten $\nu_i \in \mathbb{K}$ gibt derart, dass

$$\mathbf{w} = \sum_{i=1}^t \mu_i \mathbf{w}_i = \sum_{i=1}^r \nu_i \mathbf{d}_i,$$

also

$$\vec{0} = \sum_{i=1}^t \mu_i \mathbf{w}_i - \sum_{i=1}^r \nu_i \mathbf{d}_i,$$

und da die $\mu_1 = \dots = \mu_t = 0$ wegen der linearen Unabhängigkeit der Basisvektoren in B_W folgt auch $\nu_1 = \dots = \nu_r = 0$. Analog folgt $\lambda_1 = \dots = \lambda_s = \nu_1 = \dots = \nu_r = 0$, und somit sind die Vektoren in B_S linear unabhängig. B_S erzeugt S , und S hat demnach die Dimension $r + s + t$. Damit ist die Behauptung bewiesen. $\square$

2.4.3 Zusammenfassung für den Fall $V = \mathbb{R}^n$

1. Es sei $V_n = \mathbb{R}^n$. Jeder Vektor $\mathbf{x} = (x_1, \dots, x_n)'$ kann als Linearkombination der kanonischen Basis $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$ dargestellt werden, wobei $\mathbf{e}_j$ der j -te n -dimensionale Einheitsvektor ist:

$$\mathbf{x} = x_1 \mathbf{e}_1 + x_2 \mathbf{e}_2 + \dots + x_n \mathbf{e}_n, \quad (2.101)$$

und $V_n = \mathcal{L}(\mathbf{e}_1, \dots, \mathbf{e}_n)$. Eine Basis für V_n enthält nicht mehr als n linear unabhängige Vektoren, gleichzeitig ist es ein minimales Erzeugendensystem, da kein Basisvektor aus einer Basis entfernt werden darf, wenn *alle* Elemente von V_n erzeugbar sein sollen (s. Definition 2.20, Seite 50). Die Anzahl der Basisvektoren $\mathbf{b}_j$ ist gleich der Anzahl n der Komponenten von $\mathbf{x} \in V_n$. Diese Aussage gilt für jede Basis $\{\mathbf{b}_1, \dots, \mathbf{b}_n\}$ von V_n , da die $\mathbf{b}_k$ selbst als Linearkombinationen der $\mathbf{e}_j$ dargestellt werden können. Vergl. Satz 2.13, Seite 53.

2. Es sei $S = \{\mathbf{x}_1, \dots, \mathbf{x}_k\}$ eine Teilmenge von Vektoren aus V_n . Dann hat S den Rang r , wenn $\mathcal{L}(S)$ durch eine Basis $\{\mathbf{b}_1, \dots, \mathbf{b}_r\}$ erzeugt wird, $\mathbf{b}_j \in V_n$ für $j = 1, \dots, r$ (Definition 2.19). Für jeden Vektor $\mathbf{x} \in \mathcal{L}(S)$ existieren Koeffizienten $a_1, \dots, a_r$ derart, dass

$$\mathbf{x} = a_1 \mathbf{b}_1 + \dots + a_r \mathbf{b}_r.$$

Da die $\mathbf{b}_j$ wiederum als Linearkombinationen der kanonischen Basis repräsentiert werden können, werden letztlich alle $\mathbf{x} \in \mathcal{L}(S)$ als Linearkombinationen der kanonischen Basis dargestellt. Gleichwohl ist $\mathcal{L}(S)$ ein Teilraum des V_n .

3. Es sei wieder $S = \{\mathbf{x}_1, \dots, \mathbf{x}_k\}$ und S habe den Rang r . Dann hat das orthogonale Komplement von $\mathcal{L}(S)$ den Rang $n - r$ (vergl. Definition 2.21, Seite 55).

2.4.4 Polynome, stetige Funktionen und Vektorräume*

Bei der Betrachtung von Vektorräumen ist implizit Bezug auf die endlich-dimensionalen Vektoren $\mathbf{x} = (x_1, \dots, x_n)'$ genommen worden. Die Definition eines Vektorraums läßt aber die Betrachtung anderer Objekte als 'Vektoren' zu, – zum Beispiel Polynome vom Grad n :

$$P = a_0 + a_1x + a_2x^2 + \dots + a_nx^n \quad (2.102)$$

Im Sinne der Definition eines Vektorraums V ist ein Polynom ein Vektor. Denn ein Vektorraum ist eine Menge von Elementen, bei der die Addition zweier Elemente wieder ein Element aus V und die Multiplikation mit einem Skalar ebenfalls wieder zu einem Element aus V führt. Dementsprechend zeigt man leicht, dass Polynome in diesem Sinne Vektoren sind. Ist also

$$Q = b_0 + b_1x + b_2x^2 + \dots + b_nx^n \quad (2.103)$$

ebenfalls ein Polynom vom Grad n , so ist die Summe $P + Q$ durch

$$P + Q = a_0 + b_0 + (a_1 + b_1)x + (a_2 + b_2)x^2 + \dots + (a_n + b_n)x^n \quad (2.104)$$

ebenfalls ein Polynom vom Grad n , und für einen Skalar λ findet man, dass

$$\lambda P = \lambda a_0 + (\lambda a_1)x + \dots + (\lambda a_n)x^n \quad (2.105)$$

ebenfalls ein Polynom vom Grad n ist. Da die Koeffizienten $a_0, a_1, \dots, a_n$ auch gleich Null sein dürfen, ist die Beschränkung auf den Grad n keine Beschränkung der Allgemeinheit: man kann Polynome von höherem oder niedrigerem Grad addieren, indem man die entsprechenden Koeffizienten gleich Null setzt.

Die Elemente eines Vektorraums können als Linearkombination von Basisvektoren erzeugt werden, und so ergibt sich die Frage, welche Basisvektoren sich

für Polynome ergeben. Da die Vektoren Polynome sind, werden die Basisvektoren ebenfalls Polynome sein, und damit sie Basisvektoren sind, müssen sie linear unabhängig sein.

Satz 2.17 *Es werde der Vektorraum der Polynome vom Grad kleiner oder gleich n betrachtet. Die Polynome $1, x, x^2, \dots, x^n$ bilden eine Basis dieses Vektorraums.*

Anmerkung: die x^k sind Polynome vom Grad n , weil alle Koeffizienten a_j mit $j \neq k$ gleich Null sein dürfen.

Beweis: Zu zeigen ist, dass die $x^0 = 1, x, x^2, \dots, x^n$ linear unabhängig sind, und dass sich alle Polynome durch diese Basisvektoren erzeugen lassen. Es sei also

$$\lambda_0 + \lambda_1 x + \lambda_2 x^2 + \dots + \lambda_n x^n = 0, \quad (2.106)$$

und lineare Unabhängigkeit folgt, wenn diese Darstellung nur gilt, wenn $\lambda_0 = \lambda_1 = \dots = \lambda_n = 0$ ist. x darf Werte auf dem Intervall $[0, \infty)$ annehmen. Die λ_j dürfen nicht von x abhängen.

Es genügt, zu zeigen, dass für kein $k \geq 1$ x^k als Linearkombination der x^{k-j} , $j = 1, \dots, k-1$ dargestellt werden kann. So existieren keine Koeffizienten $\lambda_j \neq 0$, für die $x^2 = \lambda_0 + \lambda_1 x$ für alle $x \in [0, \infty)$ gelten würde. Für $x = 0$ folgt stets $\lambda_0 = 0$, und Differentiation liefert $x = \lambda_1$, entgegen der Voraussetzung, dass λ_1 unabhängig von x sein soll, also folgt $\lambda_1 = 0$. Für $\lambda_1 x + \lambda_2 x^2 + \lambda_3 x^3 = 0$ folgt nach Differentiation

$$\lambda_1 + 2\lambda_2 x + 3\lambda_3 x^2 = 0,$$

und für $x = 0$ folgt $\lambda_1 = 0$. Nochmalige Differentiation liefert $\lambda_2 + 6\lambda_3 x = 0$, so dass für $x = 0$ auch $\lambda_2 = 0$ folgt, etc. Insgesamt ist dann (2.106) nur möglich, wenn $\lambda_0 = \dots = \lambda_n = 0$, und damit folgt die Behauptung. $\square$

Der Satz 2.17 gilt für irgendein $n \in \mathbb{N}$, also für $n = 1, 2, 3, \dots$. Dann ist aber $x, x^2, x^3, \dots$ eine Basis mit unendlich vielen Elementen. Damit ist der Vektorraum der Polynome ein unendlichdimensionaler Vektorraum. Weiter ist ein Polynom als Funktion von x eine stetige Funktion, so dass die Polynome eine Teilmenge der stetigen Funktionen auf $[0, \infty]$ sind. Die Menge der auf einem Intervall – z. B. $[0, \infty)$ – stetigen Funktionen bildet ebenfalls einen Vektorraum. Denn es seien f und g irgendzwei stetige Funktionen. Dann ist die Summe $f + g$ ebenfalls eine stetige Funktion (diese Aussage läßt sich im Rahmen der Analysis beweisen), und für einen Skalar λ ist λf ebenfalls eine stetige Funktion. Die Polynome sind eine Teilmenge der stetigen Funktionen. Der Vektorraum der Polynome muß dann ein Teilvektorraum des Vektorraums der stetigen Funktionen sein, dessen Dimension (die Anzahl der Elemente der Basisvektoren) unendlich ist. Die Anzahl der Basisvektoren eines Vektorraums kann aber nicht kleiner sein als die eines Teilvektorraums. Daraus folgt, dass die Dimension des Vektorraums der stetigen Funktionen auf $[0, \infty)$ ebenfalls unendlich sein muß.

3 Matrizen

3.1 Definitionen

Wie bereits gesagt entstehen Matrizen, wenn man z.B. n jeweils m -dimensionale Vektoren $\mathbf{x}_j$, $j = 1, \dots, n$ nebeneinander schreibt, etwa

$$X = [\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n] = \begin{pmatrix} x_{11} & x_{12} & \cdots & x_{1n} \\ x_{21} & x_{22} & \cdots & x_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ x_{m1} & x_{m2} & \cdots & x_{mn} \end{pmatrix}, \quad (3.1)$$

oder indem man m n -dimensionale Vektoren als Zeilen untereinander schreibt. X ist eine " $(m \times n)$ -Matrix", m und n sind die 'Dimensionen' der Matrix. Für den Fall $m = n$ heißt X *quadratisch*. Die Zeilen einer Matrix heißen auch *Zeilenvektoren*, die Spalten *Spaltenvektoren*. Gelegentlich wird von der Schreibweise

$$X = (x_{ij}) \quad (3.2)$$

Gebrauch gemacht; damit soll gesagt werden, dass die Elemente von X mit x_{ij} bezeichnet werden.

Eine Matrix wird *gestürzt* oder *transponiert*, indem die Zeilenvektoren als Spaltenvektoren angeschrieben werden; man schreibt X' dafür:

$$X' = \begin{pmatrix} x_{11} & x_{12} & \cdots & x_{1n} \\ x_{21} & x_{22} & \cdots & x_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ x_{m1} & x_{m2} & \cdots & x_{mn} \end{pmatrix} = [\tilde{\mathbf{x}}_1, \tilde{\mathbf{x}}_2, \dots, \tilde{\mathbf{x}}_m] \quad (3.3)$$

X' ist also eine $(n \times m)$ -Matrix. Die $\tilde{\mathbf{x}}_i$, $i = 1, \dots, m$ sind die n -dimensionalen Spaltenvektoren von X' , d.h. die Zeilenvektoren von X . Ein Beispiel ist

$$X = \begin{pmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{pmatrix}, \quad X' = \begin{pmatrix} 1 & 3 & 5 \\ 2 & 4 & 6 \end{pmatrix}.$$

Wie oben schon angemerkt wurde, sind die Elemente einer Matrix nicht notwendig gemessene Daten. Die Elemente müssen also nicht notwendig reelle Zahlen sein, sie können auch komplexe Zahlen der Form $z = x + iy$ sein, wobei $x, y \in \mathbb{R}$ gilt, d.h. x und y sind reelle Zahlen, und $i = \sqrt{-1}$ ist die imaginäre Einheit. In der multivariaten Statistik hat man es zwar sehr häufig mit Matrizen mit rein reellen Elementen zu tun, aber bei der Analyse dynamischer Abläufe, also etwa bei der Zeitreihenanalyse, kann man mit komplexwertigen Elementen von Matrizen konfrontiert werden, worauf später noch explizit eingegangen wird. Der Übersicht wegen werden hier einige Typen von Matrizen vorgestellt.

Eine Matrix heißt *quadratisch*, wenn die Anzahl der Zeilen gleich der Anzahl der Spalten ist. Eine quadratische Matrix, deren Elemente oberhalb der Diagonalelemente alle gleich Null sind, heißt eine *obere Dreiecksmatrix*. Sind alle Elemente unterhalb der Diagonalelemente gleich Null, so heißt sie *untere Dreiecksmatrix*. Sind nur die Elemente in den Diagonalzellen ungleich Null, so heißt die Matrix eine *Diagonalmatrix*:

$$\Lambda = \begin{pmatrix} \lambda_1 & 0 & 0 & \cdots & 0 \\ 0 & \lambda_2 & 0 & \cdots & 0 \\ & & & \ddots & \\ 0 & 0 & 0 & \cdots & \lambda_n \end{pmatrix} = \text{diag}(\lambda_1, \dots, \lambda_n). \quad (3.4)$$

$\text{diag}(\lambda_1, \dots, \lambda_n)$ ist eine abgekürzte Schreibweise für eine Diagonalmatrix. Diagonalmatrizen werden gelegentlich auch *Skalierungsmatrizen* (engl. scaling matrix) genannt, weil man mit ihnen die Länge von Vektoren verlängern oder verkürzen, allgemein skalieren kann; worauf noch eingegangen wird. Eine Matrix heißt *symmetrisch*, wenn $x_{ij} = x_{ji}$, dh wenn das Element in der i -ten Zeile und j -ten Spalte gleich dem Element in der j -ten Zeile und der i -ten Spalte ist:

$$\begin{pmatrix} 1 & 5 \\ 5 & -2 \end{pmatrix} \quad (3.5)$$

ist ein Beispiel für eine symmetrische Matrix. Die Elemente in der Diagonalen – hier 1 und -2 – müssen also nicht identisch sein. Ist also A eine symmetrische Matrix, so gilt

$$A = A' \quad (A \text{ ist symmetrisch}). \quad (3.6)$$

Symmetrische Matrizen sind notwendig *quadratisch*, dh die Zahl der Zeilen ist stets gleich der Zahl der Spalten. Eine Diagonalmatrix ist notwendig auch symmetrisch. Es gibt zwei wichtige Spezialfälle:

1. Die Einheitsmatrix: $I = (x_{ij})$, mit $x_{ij} = \delta_{ij}$, wobei

$$\delta_{ij} = \begin{cases} 1, & i = j \\ 0, & i \neq j \end{cases}, \quad (3.7)$$

d.h.

$$I = \begin{pmatrix} 1 & 0 & 0 & \cdots & 0 \\ 0 & 1 & 0 & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \cdots & 1 \end{pmatrix} \quad (3.8)$$

Gelegentlich wird I_n geschrieben, um anzuzeigen, dass I n Zeilen und Spalten hat. Die Spalten- bzw. Zeilenvektoren sind gerade die n Einheitsvektoren $\mathbf{e}_1, \dots, \mathbf{e}_n$.

2. Die Nullmatrix: $0 = (x_{ij})$ mit $x_{ij} = 0$ für alle i und j .

Die Matrix A enthalte komplexwertige Elemente $z_{jk} = x_{jk} + iy_{jk}$, $i = \sqrt{-1}$, $x_{jk}, y_{jk} \in \mathbb{R}$. Die zu z konjugiert komplexe Zahl $\bar{z}$ ist durch $\bar{z}_{jk} = x_{jk} - iy_{jk}$

definiert. Dementsprechend heißt $\bar{A}$ die zu A *konjugiert komplexe Matrix*. Sind alle Imaginärteile y_{ij} gleich Null, so heißt die Matrix *reell* und es gilt $A = \bar{A}$, sind alle x_{ij} gleich Null und nur die Imaginärteile $y_{ij} \neq 0$, so heißt A *rein imaginär*. Die Matrix A ist rein imaginär dann und nur dann, wenn $A = -\bar{A}$.

Für die Elemente der Matrix A gelte $a_{ij} = -a_{ji}$. Dann folgt $a_{ii} = -a_{ii}$, so dass notwendig $a_{ii} = 0$. Matrizen, die diese Bedingung erfüllen, heißen *schief-symmetrisch*.

Gilt für eine Matrix $A = \bar{A}'$, ist sie also gleich der transponierten konjugiert komplexen Matrix, so heißt A *hermitesch*¹⁴, und gilt $A = -\bar{A}'$, so heißt A *schief-hermitesch*. Ist also A reell und symmetrisch, so gilt $A = \bar{A} = A'$, d.h. A ist eine reelle hermitesche Matrix, und ist A reell und schief-symmetrisch, so ist A eine reelle schief-hermitesche Matrix. Der Punkt bei diesen Definitionen ist, dass alle Eigenschaften von allgemeinen komplexen hermiteschen Matrizen auch für den Spezialfall reeller hermitescher Matrizen gelten, da die reellen Zahlen ja Spezialfälle komplexer Zahlen sind. Reelle symmetrische Matrizen spielen in der multivariaten Statistik eine zentrale Rolle (etwa die Varianz-Kovarianz-Matrizen).

Beispiel 3.1 Matrix von Skalarprodukten Gegeben seien die Vektoren

$$\mathbf{x}_1, \dots, \mathbf{x}_n$$

und es werden die Skalarprodukte $\mathbf{x}_j' \mathbf{x}_k$ für $j, k = 1, \dots, n$ berechnet. Die Skalarprodukte lassen sich in einer Matrix anordnen: die erste Zeile enthalte die Skalarprodukte $\mathbf{x}_1' \mathbf{x}_k$, $k = 1, \dots, n$, die zweite Zeile enthalte die Skalarprodukte $\mathbf{x}_2' \mathbf{x}_k$ für $k = 1, \dots, n$, etc. Es entsteht eine symmetrische Matrix mit den Elementen $x_{jk} = \mathbf{x}_j' \mathbf{x}_k$. Bei entsprechender Normierung sind die Skalarprodukte Kovarianzen bzw. Korrelationen, – Kovarianz- und Korrelationsmatrizen sind symmetrisch. Sind die $\mathbf{x}_j$ paarweise orthogonal, so entsteht eine Diagonalmatrix. □

Submatrizen: Eine Matrix kann in Teil- oder Submatrizen aufgeteilt werden, bzw. aus solchen Matrizen aufgebaut werden. Es sei A eine $m \times n$ -Matrix. A kann z.B. aus Teilmatrizen A_{11} , A_{12} , A_{21} , A_{22} zusammengesetzt sein:

$$A = \begin{pmatrix} A_{11} & A_{12} \\ A_{21} & A_{22} \end{pmatrix} \quad (3.9)$$

wobei A_{11} eine $r \times s$ -Matrix, A_{12} eine $r \times (n - s)$ -Matrix, A_{21} eine $(m - r) \times s$ -Matrix, A_{22} eine $(m - r) \times (n - s)$ -Matrix ist. Diese Aufteilung läßt sich leicht verallgemeinern auf $p \times q$ Teilmatrizen, wenn nur m und n hinreichend groß sind.

Für die Transponierte A' hat man

$$A' = \begin{pmatrix} A'_{11} & A'_{12} \\ A'_{21} & A'_{22} \end{pmatrix} \quad (3.10)$$

¹⁴Charles Hermite (1822 – 1901), französischer Mathematiker.

Die Spur einer quadratischen Matrix: Es sei A eine $n \times n$ -Matrix, d.h. A habe so viele Zeilen wie Spalten. A heißt dann *quadratisch*. Die Summe $\text{spur}(A)$ der Diagonalelemente heißt dann die *Spur* der Matrix A :

$$\text{spur}(A) = \sum_{i=1}^n a_{ii}. \quad (3.11)$$

Der Begriff der Spur einer Matrix verweist nicht unmittelbar auf eine bereits bekannte Operation mit Vektoren. Der Sinn dieses Begriffs ergibt sich erst in späteren Anwendungen der Matrixrechnung.

3.2 Elementare Operationen mit Matrizen

Es sei $X = (x_{ij})$, $i = 1, \dots, m$, $j = 1, \dots, n$, d.h. X sei eine $(m \times n)$ -Matrix mit den Elementen x_{ij} , wobei i für die i -te Zeile und j für die j -te Spalte steht.

λ sei ein Skalar (d.h. eine reelle Zahl). Dann gilt

$$\lambda X = (\lambda x_{ij}), \quad (3.12)$$

d.h. die Multiplikation von X mit λ bedeutet, dass jedes Element x_{ij} von X mit λ multipliziert wird.

Es seien X und Y zwei Matrizen mit identischen Dimensionen, d.h. beide Matrizen seien $(m \times n)$ -Matrizen. Dann gilt

$$X + Y = (x_{ij} + y_{ij}), \quad (3.13)$$

D.h. die Matrizen werden addiert, indem man die zueinander korrespondierenden Elemente addiert. Weiter gilt

$$(X + Y)' = X' + Y'. \quad (3.14)$$

Ist $\lambda \in \mathbb{R}$, so sieht man sofort, dass

$$\text{spur}(\lambda A) = \lambda \text{spur}(A). \quad (3.15)$$

Ist B ebenfalls eine $n \times n$ -Matrix, so folgt aus der Definition der Addition von Matrizen ebenfalls sofort

$$\text{spur}(A + B) = \text{spur}(A) + \text{spur}(B). \quad (3.16)$$

Es sei nun C eine $n \times p$ -Matrix und D eine $p \times n$ -Matrix. Es sei $A = CD$; A ist eine quadratische, d.h. eine $n \times n$ -Matrix. Für das Element a_{ii} hat man dann

$$a_{ii} = \sum_{j=1}^p b_{ij} c_{ji}.$$

Dann folgt

$$\text{spur}(A) = \text{spur}(CD) = \sum_{i=1}^n \sum_{j=1}^p b_{ij} c_{ji}. \quad (3.17)$$

Ein Spezialfall von (3.17) ergibt sich für $C = D$:

$$\text{spur}(CC') = \text{spur}(C'C) = \sum_{i=1}^n \sum_{j=1}^n c_{ij}^2. \quad (3.18)$$

3.3 Die Multiplikation von Matrizen

3.3.1 Die Multiplikation einer Matrix mit einem Vektor

Gegeben sei eine Linearkombination von n m -dimensionalen Vektoren $\mathbf{x}_j$, $j = 1, \dots, n$:

$$\mathbf{y} = u_1 \mathbf{x}_1 + u_2 \mathbf{x}_2 + \dots + u_n \mathbf{x}_n \quad (3.19)$$

Die $\mathbf{x}_j$ können zu einer $(m \times n)$ -Matrix X zusammengefasst werden, und die $u_1, \dots, u_n$ können als Komponenten eines Vektors $\mathbf{u}$ aufgefasst werden. Für die rechte Seite von (3.19) werde dann die Schreibweise $X\mathbf{u}$ eingeführt, so dass

$$\mathbf{y} = X\mathbf{u} =_{\text{def}} u_1 \mathbf{x}_1 + u_2 \mathbf{x}_2 + \dots + u_n \mathbf{x}_n \quad (3.20)$$

resultiert. Das *Produkt $X\mathbf{u}$ einer Matrix mit einem Vektor* soll also einfach für eine Linearkombination der Spaltenvektoren der Matrix stehen, wobei die Koeffizienten u_i , $i = 1, \dots, n$ der Spaltenvektoren durch die Komponenten des Vektors $\mathbf{u}$ gegeben sind.

Gleichzeitig bietet die Schreibweise $\mathbf{y} = X\mathbf{u}$ eine weitere Interpretation an, nämlich die der *Transformation* des n -dimensionalen Vektors $\mathbf{u}$ in den m -dimensionalen Vektor $\mathbf{y}$, bzw. die einer *Abbildung* des Vektors $\mathbf{u}$ auf den Vektor $\mathbf{y}$, und dementsprechend kann man sagen, dass die Matrix X diese Transformation oder Abbildung definiert:

$$X: \mathbf{u} \mapsto \mathbf{y}. \quad (3.21)$$

Auf diese Interpretation wird in Abschnitt 3.3.3 näher eingegangen.

Schreibt man die Linearkombination in (3.19) aus, so erhält man

$$\mathbf{y} = X\mathbf{u} = \begin{pmatrix} x_{11}u_1 + x_{12}u_2 + \dots + x_{1n}u_n \\ x_{21}u_1 + x_{22}u_2 + \dots + x_{2n}u_n \\ \vdots \\ x_{m1}u_1 + x_{m2}u_2 + \dots + x_{mn}u_n \end{pmatrix} \quad (3.22)$$

Dass die rechte Seite in die Form

$$\mathbf{y} = u_1 \begin{pmatrix} x_{11} \\ x_{21} \\ \vdots \\ x_{m1} \end{pmatrix} + u_2 \begin{pmatrix} x_{12} \\ x_{22} \\ \vdots \\ x_{m2} \end{pmatrix} + \dots + u_n \begin{pmatrix} x_{1n} \\ x_{2n} \\ \vdots \\ x_{mn} \end{pmatrix}$$

gebracht werden kann, sieht man leicht. Wichtig ist aber noch ein anderer Aspekt von (3.22): die Zeilen auf der rechten Seite sind einerseits die Komponenten des Vektors $\mathbf{y}$, andererseits haben sie die Form von Skalarprodukten. Für die i -te Komponente y_i von $\mathbf{y}$ hat man

$$y_i = x_{i1}u_1 + x_{i2}u_2 + \cdots + x_{in}u_n, \quad (3.23)$$

d.h. y_i ist gleich dem Skalarprodukt des i -ten Zeilenvektors von X mit dem Vektor $\mathbf{u}$. Bezeichnet man also den i -ten Zeilenvektor von X mit $\tilde{\mathbf{x}}_i$, so kann man für $\mathbf{y} = X\mathbf{u}$ auch

$$\mathbf{y} = \begin{pmatrix} \tilde{\mathbf{x}}_1 \cdot \mathbf{u} \\ \tilde{\mathbf{x}}_2 \cdot \mathbf{u} \\ \vdots \\ \tilde{\mathbf{x}}_m \cdot \mathbf{u} \end{pmatrix} \quad (3.24)$$

schreiben (hier wurde von der Schreibweise $\tilde{\mathbf{x}}_1 \cdot \mathbf{u}$ statt von $\tilde{\mathbf{x}}_1' \mathbf{u}$ Gebrauch gemacht, da diese voraussetzt, dass $\tilde{\mathbf{x}}_i$ als Spaltenvektor aufgefasst wird, was zu Konfusionen führen könnte, da die $\tilde{\mathbf{x}}_i$ ja schon als Zeilenvektoren eingeführt wurden). Bei der tatsächlichen Berechnung von $\mathbf{y} = X\mathbf{u}$ wird im Allgemeinen von (3.24) Gebrauch gemacht. Diese Interpretation von $X\mathbf{u}$ erweist sich auch als nützlich, wenn die Komponenten x_{ij} und u_j Abweichungen von Mittelwerten sind, da die $y_i = \tilde{\mathbf{x}}_i \cdot \mathbf{u}$ dann zu Kovarianzen korrespondieren.

Man rechnet leicht nach, dass

$$(X\mathbf{u})' = \mathbf{u}'X' = \mathbf{y}'. \quad (3.25)$$

Hier wird eine Matrix $M = X'$ *von links* mit einem Zeilenvektor $\mathbf{u}'$ multipliziert, und es entsteht offenbar ein Zeilenvektor, nämlich $\mathbf{y}'$. Nun war aber $\mathbf{y}$ eine Linearkombination der Spalten von X , die die Zeilen von X' sind. Dies legt nahe, zu sagen, dass nach (3.25) der Zeilenvektor $\mathbf{y}'$ eine Linearkombination der Zeilenvektoren von $M (= X')$ ist. Dass diese Aussage tatsächlich generell gilt, wird im Folgenden elaboriert.

Es sei also $\mathbf{u}$ ein m -dimensionaler Vektor, und es werden die Skalarprodukte des Zeilenvektors $\mathbf{u}'$ mit den Spaltenvektoren von X gebildet. Dann ist

$$\begin{aligned} \mathbf{z}' &= \mathbf{u}'X = \left(\sum_{i=1}^m u_i x_{i1}, \sum_{i=1}^m u_i x_{i2}, \dots, \sum_{i=1}^m u_i x_{in} \right) \\ &= u_1(x_{11}, x_{12}, \dots, x_{1n}) + \cdots + u_n(x_{m1}, x_{m2}, \dots, x_{mn}) \end{aligned} \quad (3.26)$$

Das Produkt $\mathbf{u}'X$ liefert also tatsächlich einen Zeilenvektor. Zusammenfassend hat man also:

1. **Multiplikation von rechts:** Wird eine Matrix X *von rechts* mit einem *Spaltenvektor* $\mathbf{u}$ multipliziert, so entsteht ein Spaltenvektor $\mathbf{y}$ als Linearkombination der Spaltenvektoren $\mathbf{x}_1, \dots, \mathbf{x}_n$ von X .

Die Komponenten von $\mathbf{y}$ sind die Skalarprodukte der *Zeilenvektoren* von X mit dem Vektor $\mathbf{u}$.

2. **Multiplikation von links** Wird eine Matrix *von links* mit einem Zeilenvektor $\mathbf{u}'$ multipliziert, so entsteht ein Zeilenvektor $\mathbf{z}'$ als Linearkombination der *Zeilenvektoren* von X .

Die Komponenten von $\mathbf{z}'$ sind die Skalarprodukte des Zeilenvektors $\mathbf{u}'$ mit den *Spaltenvektoren* von X .

Skalierung von Vektoren Eine spezielle Transformation von Vektoren ist die Skalierung ihrer Länge. Diese kann durch die Multiplikation mit Diagonalmatrizen erreicht werden, die deswegen auch Skalierungsmatrizen (scaling matrices) genannt werden (vergl. Seite 62). Dieser Sachverhalt kann anhand von (2×2) -Diagonalmatrizen illustriert werden: es ist

$$\Lambda \mathbf{x} = \begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix} \begin{pmatrix} x_{11} & x_{12} & x_{13} \\ x_{21} & x_{22} & x_{23} \end{pmatrix} = \begin{pmatrix} \lambda_1 x_{11} & \lambda_1 x_{12} & \lambda_1 x_{13} \\ \lambda_2 x_{21} & \lambda_2 x_{22} & \lambda_2 x_{23} \end{pmatrix}, \quad (3.27)$$

d.h. die Multiplikation einer $(m \times n)$ -Matrix X *von links* mit einer $(m \times m)$ -Diagonalmatrix Λ bedeutet eine Skalierung der Länge der Zeilenvektoren von X . Der i -te Zeilenvektor wird mit dem i -ten Diagonalelement von Λ skaliert.

Die Multiplikation von X *von rechts* mit einer Diagonalmatrix bedeutet die Skalierung der Spaltenvektoren von X , wobei die Diagonalmatrix aber eine $(n \times n)$ -Matrix sein muß:

$$\begin{pmatrix} x_{11} & x_{12} & x_{13} \\ x_{21} & x_{22} & x_{23} \end{pmatrix} \begin{pmatrix} \lambda_1 & 0 & 0 \\ 0 & \lambda_2 & 0 \\ 0 & 0 & \lambda_3 \end{pmatrix} = \begin{pmatrix} \lambda_1 x_{11} & \lambda_2 x_{12} & \lambda_3 x_{13} \\ \lambda_1 x_{21} & \lambda_2 x_{22} & \lambda_3 x_{23} \end{pmatrix} \quad (3.28)$$

Die Länge des j -ten Spaltenvektors von X wird mit dem j -ten Diagonalelement von Λ skaliert.

3.3.2 Der allgemeine Fall

Die Multiplikation einer Matrix mit einer Matrix ist einfach eine Verallgemeinerung der Multiplikation einer Matrix mit einem Vektor. Es sei A eine $(m \times n)$ -Matrix und B sei eine $(n \times r)$ -Matrix. Die Spalten von A seien die Vektoren $\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n$, und die Spalten von B seien die Vektoren $\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_r$, so dass man B in der Form $B = [\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_r]$ schreiben kann. Dann ist das Produkt AB definiert durch

$$AB = [A\mathbf{b}_1, A\mathbf{b}_2, \dots, A\mathbf{b}_r] = [\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_r] = C, \quad (3.29)$$

wobei die $A\mathbf{b}_j = \mathbf{c}_j$ offenbar m -dimensionale Vektoren sind, d.h. C ist eine $(m \times r)$ -Matrix. $\mathbf{c}_j$ ist eine Linearkombination der Spaltenvektoren von A , $j = 1, \dots, r$, und die Komponenten von $\mathbf{b}_j$ sind die entsprechenden Koeffizienten.

Nun sei α_i die i -te Zeile bzw. der i -te Zeilenvektor von A . Dann ist

$$\alpha_i B = [\alpha_i \mathbf{b}_1, \alpha_i \mathbf{b}_2, \dots, \alpha_i \mathbf{b}_r] = (c_{i1}, c_{i2}, \dots, c_{ir}). \quad (3.30)$$

vergl. Punkt 2, Seite 67. d.h. $\alpha_i B$ ist der i -te Zeilenvektor von C (in (3.30) wurde nicht $\alpha'_i B$ geschrieben, weil α_i bereits als Zeilenvektor eingeführt wurde), und die als Vektor aufgefasste Zeile $(c_{i1}, c_{i2}, \dots, c_{ir})$ ist eine Linearkombination der Zeilenvektoren von B .

Die Elemente c_{ij} von C sind Skalarprodukte der Zeilenvektoren von A mit den Spaltenvektoren von B :

$$\begin{aligned} AB &= \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix} \begin{pmatrix} b_{11} & b_{12} & \cdots & b_{1r} \\ b_{21} & b_{22} & \cdots & b_{2r} \\ \vdots & \vdots & \ddots & \vdots \\ b_{n1} & b_{n2} & \cdots & b_{nr} \end{pmatrix} \\ &= \begin{pmatrix} \sum_{j=1}^n a_{1j} b_{j1} & \sum_{j=1}^n a_{1j} b_{j2} & \cdots & \sum_{j=1}^n a_{1j} b_{jr} \\ \vdots & \vdots & \ddots & \vdots \\ \sum_{j=1}^n a_{mj} b_{j1} & \sum_{j=1}^n a_{mj} b_{j2} & \cdots & \sum_{j=1}^n a_{mj} b_{jr} \end{pmatrix} = C \quad (3.31) \end{aligned}$$

Notwendige Voraussetzung für die Berechnung eines Produkts zweier Matrizen A und B ist, dass die Anzahl der Spalten von A gleich der Anzahl von Zeilen von B ist.

Matrixmultiplikation und Linearkombinationen: Das Ergebnis sei noch einmal zusammengefasst: Es sei $C = AB$. Dann gilt

1. Die *Spaltenvektoren* von C sind Linearkombinationen der *Spaltenvektoren* von A . Dies folgt aus den Betrachtungen zum Produkt $\mathbf{y} = X\mathbf{u}$, Gleichung (3.20), Seite 65.
2. Die *Zeilenvektoren* von C sind Linearkombinationen der *Zeilenvektoren* von B . Dies folgt aus den Betrachtungen zu (3.26), Seite 66.

Für die Matrixmultiplikation gelten die folgenden Aussagen, die sich aus den beiden Regeln 1 und 2 herleiten lassen:

$$AB \neq BA \quad (3.32)$$

$$(AB)' = B'A' \quad (3.33)$$

$$(AB)C = A(BC) \quad (3.34)$$

(3.32) bedeutet, dass die Matrixmultiplikation *nicht kommutativ* ist, – die Multiplikation von Skalaren ist dagegen kommutativ: $ab = ba$, also $3 \times 4 = 4 \times 3$. Das Produkt AB bedeutet, dass die Spaltenvektoren der entstehenden Matrix C Linearkombinationen der Spalten von A sind, während BA bedeutet, dass die

Spaltenvektoren der entstehenden Matrix D Linearkombinationen der Spaltenvektoren von B sind, so dass im Allgemeinen $AB \neq BA$ gilt. Darüber hinaus muß für die Bildung des Produkts AB die Anzahl der Spalten von A gleich der Anzahl der Zeilen von B sein, und für BA muß die Anzahl der Spalten von B gleich der Anzahl der Zeilen von A sein, aber diese Korrespondenz von Spalten- und Zeilenanzahlen muß keineswegs gegeben sein.

Die Gleichung (3.33) ergibt sich ebenfalls aus den Punkten 1 und 2. Es sei $C = AB$; die Spalten von C sind Linearkombinationen der Spalten von A , die Zeilen von C sind Linearkombinationen der Zeilen von B . Dann müssen die Spalten von C' Linearkombinationen der Spalten von B' sein, und hieraus folgt sofort $(AB)' = B'A'$.

Die Gleichung (3.34) bedeutet die *Assoziativität* der Matrixmultiplikation. Der Ausdruck $(AB)C$ bedeutet ja, dass das Produkt die Linearkombinationen der Spalten von AB enthält, und AB wiederum besteht aus den Linearkombinationen der Spalten von A . Damit bestehen die Spalten von $(AB)C$ letztlich aus Linearkombinationen der Spalten von A . Aber genau dies wird mit dem Produkt $A(BC)$ ausgedrückt.

Zur Übung mache man sich die Richtigkeit der Aussagen an einem kleinen numerischen Beispiel klar.

Anmerkung: Die Definition des dyadischen Produkts (Def. 2.6, Seite 29) ergibt sich aus der für die Matrixmultiplikation, wenn man zu $\mathbf{x}$ eine Anzahl von Nullvektoren $\vec{0}$ hinzufügt und zu $\mathbf{y}'$ eine entsprechende Zahl von Zeilenvektoren $\vec{0}'$:

$$\begin{pmatrix} x_1 \\ x_2 \end{pmatrix} (y_1, y_2, y_3) = \begin{pmatrix} x_1 & 0 \\ x_2 & 0 \end{pmatrix} \begin{pmatrix} y_1 & y_2 & y_3 \\ 0 & 0 & 0 \end{pmatrix} = \begin{pmatrix} x_1 y_1 & x_1 y_2 & x_1 y_3 \\ x_2 y_1 & x_2 y_2 & x_2 y_3 \end{pmatrix}$$

Es gibt spezielle Matrizen, die insbesondere in der multivariaten Statistik gelegentlich vorkommen

Definition 3.1 *Es sei M eine quadratische Matrix. M heißt idempotent, wenn*

$$MM = M^2 = M, \quad (3.35)$$

d.h. das Produkt von M mit sich selbst liefert wieder M .

Im nächsten Abschnitt wird ein Beispiel geliefert.

3.3.3 Transformationen und Abbildungen

Dieser Abschnitt enthält einige grundsätzliche Betrachtungen über Produkte von Matrizen und Vektoren, die einerseits das Verständnis der Vektor- und Matrixrechnung vertiefen, andererseits für das Verständnis der unmittelbaren Anwendung

der Matrixrechnung auf Fragen der multivariaten Statistik nicht unbedingt notwendig sind und deshalb übersprungen werden können.

Das Produkt $X\mathbf{u} = \mathbf{v}$, X eine $(m \times n)$ -Matrix, $\mathbf{u}$ ein n -dimensionaler Vektor, $\mathbf{v}$ ein m -dimensionaler Vektor kann als Abbildung $f: \mathbf{u} \mapsto \mathbf{v}$ eines n -dimensionalen Vektors auf einen m -dimensionalen Vektor verstanden werden, wobei die Abbildung f durch die Matrix X definiert wird. Das Gleiche gilt für das Produkt $\mathbf{u}'X = \mathbf{v}'$, wenn $\mathbf{u}$ ein m -dimensionaler und $\mathbf{v}$ ein n -dimensionaler Vektor ist. Viele Sachverhalte der Vektor- und Matrixalgebra lassen sich sehr elegant als Eigenschaften von Abbildungen ausdrücken. In Abschnitt 4 werden Abbildungen ausführlicher diskutiert, hier werden nur einige wesentliche Aspekte vorgestellt.

Eine Abbildung f einer Menge $\mathcal{M}$ in eine Menge $\mathcal{N}$ ordnet jedem Element aus $\mathcal{M}$ *genau einem* Element aus $\mathcal{N}$ zu:

$$f: \mathcal{M} \rightarrow \mathcal{N}, \quad x \mapsto y = f(x), \quad x \in \mathcal{M}, y \in \mathcal{N} \quad (3.36)$$

Man schreibt gelegentlich auch

$$f(\mathcal{M}) = \mathcal{N}. \quad (3.37)$$

$f(\mathcal{M})$ heißt das *Bild* von $\mathcal{M}$ in $\mathcal{N}$, und $\mathcal{M}$ ist das *Urbild* von $f(\mathcal{M})$. Man schreibt auch $\text{Im}f = \mathcal{N}$.

Eine spezielle Abbildung ist die *Identität* oder identische Abbildung

$$\text{id}(\mathcal{M}) = \mathcal{M}. \quad (3.38)$$

Die Einheitsmatrix I_n der Spalten bzw. Zeilen aus den n -dimensionalen Einheitsvektoren $\mathbf{e}_1, \dots, \mathbf{e}_n$ bestehen, spezifiziert die identische Abbildung, denn sicherlich gilt

$$I_n \mathbf{x} = \mathbf{x}, \quad \mathbf{x} \in \mathbb{R}^n. \quad (3.39)$$

Für eine Teilmenge von Abbildungen existiert die *inverse Abbildung* f^{-1} :

$$f(\mathcal{M}) = \mathcal{N}, \quad f^{-1}f(\mathcal{M}) = \mathcal{M} = f^{-1}(\mathcal{N}). \quad (3.40)$$

Wenn f durch eine Matrix M definiert ist, so bedeutet die Existenz der inversen Abbildung f^{-1} die Existenz einer inversen Matrix M^{-1} . Es wird deutlich werden, dass inverse Matrizen M^{-1} für eine Matrix M nur für spezielle Matrizen existieren.

Die Forderung, dass einem Element $x \in \mathcal{M}$ nur *ein* Element $y \in \mathcal{N}$ zugeordnet wird schließt nicht aus, dass verschiedenen Elementen $x, x' \in \mathcal{M}$ der gleiche Wert $y \in \mathcal{N}$ zugeordnet werden kann. In diesem Fall kann von einem Element $y \in \mathcal{N}$ nicht eindeutig auf das Element $x \in \mathcal{M}$ mit $f(x) = y$ zurückgeschlossen werden.

Mit der Schreibweise $f(\mathcal{M})$ ist nicht ein einzelnes Element gemeint, sondern die Menge der Werte, die man erhält, wenn man f für alle Werte aus X bestimmt, also

$$f(\mathcal{M}) = \{f(x), x \in \mathcal{M}\}. \quad (3.41)$$

Offenbar gilt $f(\mathcal{M}) \subseteq \mathcal{N}$.

Definition 3.2 Es sei $f: \mathcal{M} \rightarrow \mathcal{N}$. Dann ist f

1. injektiv, wenn aus $x, x' \in \mathcal{M}$ und $f(x) = f(x')$ folgt, dass $x = x'$ (und damit $f(x) \neq f(x') \Rightarrow x \neq x'$). Es kann $f(\mathcal{M}) \subset \mathcal{N}$ gelten, d.h. $f(\mathcal{M})$ kann eine echte Teilmenge von $\mathcal{N}$ sein.
2. surjektiv, wenn $f(\mathcal{M}) = \mathcal{N}$, d.h. zu jedem $y \in \mathcal{N}$ existiert ein $x \in \mathcal{M}$ derart, dass $y = f(x)$. Es gilt $f(\mathcal{M}) = \mathcal{N}$.
3. bijektiv, wenn f sowohl injektiv als auch surjektiv ist. Es gilt $f(\mathcal{M}) = \mathcal{N}$.
4. Die Menge $\text{kern} f = \{x \in \mathcal{M} | f(x) = 0\}$ heißt Kern der Abbildung f ; man schreibt für den Kern auch $\text{kern} f = f^{-1}(\vec{0})$.
5. Es sei $f(\mathcal{M}) = \mathcal{N}$, d.h. $f(\mathbf{x}) = \mathbf{y}$ für $\mathbf{x} \in \mathcal{M}$, $\mathbf{y} \in \mathcal{N}$. Dann heißt $f^{-1}(\mathbf{y}) = \{\mathbf{x} \in \mathcal{M} | f(\mathbf{x}) = \mathbf{y}\}$ die Faser über $\mathbf{y} \in \mathcal{N}$.

Anmerkung: Die Schreibweise $f^{-1}(f)$ für den Kern einer Abbildung f ergibt sich aus der in 4. gegebenen Definition: ist $\mathbf{x} \in \text{kern} f$, so gilt $f(\mathbf{x}) = \vec{0}$. Aus der Definition der Inversen folgt dann $\mathbf{x} = f^{-1}(\vec{0})$. Die Definition des Kerns setzt wie die Definition der Faser offenbar voraus, dass die Inverse existiert. $\square$

Beispiele: $f: \mathbb{R} \rightarrow \mathbb{R}$, $x \mapsto ax + b$ für $a, b \in \mathbb{R}$ fest gewählte Konstante. f ist sicher injektiv, denn $f(x) = f(x')$ impliziert $ax + b = ax' + b$ und damit $x = x'$, wie man leicht nachrechnet. f ist auch surjektiv, denn für $y = ax + b$ existiert genau ein $x = (y - b)/a$ derart, dass $y = f(x)$. Da f sowohl injektiv wie surjektiv ist, ist f auch bijektiv.

$\mathbb{R}$ bezeichnet die Menge der reellen Zahlen. Mit $\mathbb{R} \times \mathbb{R} = \mathbb{R}^2$ wird die Menge der Paare (x, y) , $x, y \in \mathbb{R}$, bezeichnet, allgemein mit

$$\mathbb{R}^m = \underbrace{\mathbb{R} \times \mathbb{R} \times \cdots \times \mathbb{R}}_{m\text{-mal}}$$

Die Menge der m -tupel $(x_1, x_2, \dots, x_m)$, $x_j \in \mathbb{R}$, d.h. der m -dimensionalen Vektoren. $\mathbb{R}^n$ ist dann die Menge der n -dimensionalen Vektoren, etc. Mit $\mathbb{R}^{m,n}$ wird die Menge der $(m \times n)$ -Matrizen bezeichnet. Alle diese Definitionen übertragen sich auf $\mathbb{C}$, die Menge der komplexen Zahlen $x + iy$, $x, y \in \mathbb{R}$ und $i = \sqrt{-1}$.

Die Schreibweise $M \in \mathbb{R}^{m,n}$ bedeutet, dass M eine $(m \times n)$ -Matrix ist. Die Schreibweise $f: V_m \rightarrow V_n$ bedeutet dann, dass f eine Abbildung der m -dimensionalen Vektoren in die Menge der n -dimensionalen Vektoren ist. Wenn $V_m = \mathbb{R}^m$, $V_n = \mathbb{R}^n$ kann man auch $f: \mathbb{R}^m \rightarrow \mathbb{R}^n$ schreiben.

Da $M\mathbf{x} = \mathbf{y}$ mit $\mathbf{x} \in V_m$, $\mathbf{y} \in V_n$, folgt, dass f durch eine Matrix $M \in \mathbb{R}^{m,n}$ definiert ist.

Definition 3.3 Es seien V und W Vektorräume und $f: V \rightarrow W$ sei eine Abbildung von V in W . f heißt linear bzw. homomorph, wenn

$$f(\lambda \mathbf{v} + \mu \mathbf{w}) = \lambda f(\mathbf{v}) + \mu f(\mathbf{w}) \quad (3.42)$$

für $\lambda, \mu \in \mathbb{R}$ und für alle $\mathbf{v} \in V$ und $\mathbf{w} \in W$. Insbesondere heißt f isomorph, wenn f bijektiv ist; man sagt auch, f definiere einen Homomorphismus bzw. Isomorphismus für f bijektiv. f definiert einen Endomorphismus, wenn $V = W$, und einen Automorphismus, wenn f bijektiv ist und außerdem $V = W$ gilt.

Es sei $M \in \mathbb{R}^{m,n}$; M definiert eine lineare, also homomorphe Abbildung, denn $M\mathbf{x} = \mathbf{y}$ erfüllt die Bedingungen einer linearen Abbildung. für $m \neq n$ ist f offenbar weder ein Endomorphismus noch ein Automorphismus.

Dem Begriff des Kerns in der allgemeinen Definition 3.2 von Abbildungen entspricht für $f \in \mathbb{R}^{m,n}$ der Nullvektor $\vec{0}$.

Satz 3.1 Es sei $f(\mathcal{M}) = \mathcal{N}$. Dann gilt

1. f ist surjektiv genau dann, wenn $\text{Im } f = \mathcal{N}$
2. f ist injektiv genau dann, wenn $\text{kern } f = \{\vec{0}\}$.
3. f sei injektiv und die Vektoren $\mathbf{x}_1, \dots, \mathbf{x}_n \in \mathcal{M}$ seien linear unabhängig. Dann sind auch die Bilder $f(\mathbf{x}_1), \dots, f(\mathbf{x}_n)$ linear unabhängig.

Anmerkung: Die Schreibweise $\text{kern } f = \{\vec{0}\}$ bedeutet, dass $\text{kern } f$ nur das eine Element $\vec{0}$ enthält. $\square$

Beweis: $\Rightarrow$ für 1. und 2. folgt sofort aus der Definition von injektiv und surjektiv. Um $\Leftarrow$ zu sehen, betrachte man zwei Vektoren $\mathbf{u}, \mathbf{v} \in \mathcal{M}$ mit $\mathbf{u} \neq \mathbf{v}$, aber $f(\mathbf{u}) = f(\mathbf{v})$. Wegen der Linearität von f folgt dann

$$f(\mathbf{v}) - f(\mathbf{u}) = f(\mathbf{v} - \mathbf{u}) = \vec{0},$$

d.h. es gilt $\mathbf{v} - \mathbf{u} \in \text{kern } f$.

Um 3. einzusehen sei angenommen, dass

$$\lambda_f(\mathbf{x}_1) + \dots + \lambda_n f(\mathbf{x}_n) = \vec{0}$$

gilt. Es wurde vorausgesetzt, dass f injektiv ist. Daraus folgt, dass

$$\lambda_1 \mathbf{x}_1 + \dots + \lambda_n \mathbf{x}_n = \vec{0}$$

gelten muß, denn $\lambda_1 \mathbf{x}_1 + \dots + \lambda_n \mathbf{x}_n$ ist ja das Urbild von f . Da die $\mathbf{x}_1, \dots, \mathbf{x}_n$ als linear unabhängig vorausgesetzt wurden, muß $\lambda_1 = \dots = \lambda_n = 0$ gelten, und dann folgt sofort, dass auch die $f(\mathbf{x}_j)$ linear unabhängig sind. $\square$

Definition 3.4 Es sei f eine Abbildung $f: \mathbb{R}^m \rightarrow \mathbb{R}^n$; dann heißt die Dimensionalität des Bildes $\text{Im } f$ der Rang; man schreibt $\text{rg}(f) = \dim \in f$.

f sei durch eine Matrix $A \in \mathbb{R}^{m,n}$ definiert, so dass $A: \mathbb{R}^m \rightarrow \mathbb{R}^n$ und $\mathbf{x} \mapsto \mathbf{y} = A\mathbf{x}$. Dann ist $f = A \cdot \mathbf{e}_1, \dots, \mathbf{e}_m$ ist die kanonische Basis von $\mathbb{R}^m$, d.h. die n m -dimensionalen Spaltenvektoren von A können als Linearkombinationen

$A\mathbf{e}_1, A\mathbf{e}_2, \dots, A\mathbf{e}_n$ geschrieben werden. Dann ist das Bild der durch A definierten Abbildung die lineare Hülle

$$\operatorname{Im} A = \mathcal{L}(A\mathbf{e}_1, A\mathbf{e}_2, \dots, A\mathbf{e}_n).$$

Demnach wird $\operatorname{Im} A$ auch der *Spaltenraum* von A bezeichnet. Der Begriff des Ranges einer Matrix A wird in Abschnitt 3.6 noch ausführlich diskutiert.

Beispiel 3.2 Es sei $f: \mathbb{R}^3 \rightarrow \mathbb{R}^4$; für $\mathbf{x} \in \mathbb{R}^3$ und $\mathbf{y} \in \mathbb{R}^4$ soll also $f(\mathbf{x}) = \mathbf{y}$ gelten; insbesondere sei f durch

$$f(\mathbf{x}) = \begin{pmatrix} x_1 + 2x_2 \\ 0 \\ \frac{1}{2}x_1 + x_2 \\ 0 \end{pmatrix}$$

definiert. Gesucht ist die zu f gehörige Matrix $M = A$ sowie der Kern von f .

Der Kern von f ist diejenige Menge von Vektoren $\mathbf{x}$, für die $f(\mathbf{x}) = \vec{0}$. Für dies Komponenten dieser Vektoren $\mathbf{x}$ muß also gelten

$$\begin{aligned} x_1 + 2x_2 &= 0 \\ \frac{1}{2}x_1 + x_2 &= 0, \end{aligned}$$

d.h. $x_1 = -2x_2$. Der Kern ist dann

$$\operatorname{kern}(f) = \{(x_1, x_2, x_3)' | x_1 = -2x_2\} = \{(-2x_2, x_2, x_3)'\}.$$

Es gilt

$$A\mathbf{x} = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \\ a_{41} & a_{42} & a_{43} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \end{pmatrix} = \mathbf{y}.$$

Es gibt also 12 Elemente a_{ij} , die zu bestimmen sind, wobei allerdings nur bestimmte Relationen zwischen den Komponenten gegeben sind, die aus dem Spezialfall $A\mathbf{x} = \vec{0}$ folgen. Wie die Diskussion linearer Gleichungssysteme zeigen wird, läßt sich aus diesen Bedingungen keine eindeutige Lösung für die a_{ij} ableiten.

Andererseits ist das Bild von f eine Linearkombination der Spalten von A , und damit folgt

$$\begin{pmatrix} x_1 + 2x_2 \\ 0 \\ \frac{1}{2}x_1 + x_2 \\ 0 \end{pmatrix} = \frac{x_1}{2} \begin{pmatrix} 2 \\ 0 \\ 1 \\ 0 \end{pmatrix} + x_2 \begin{pmatrix} 2 \\ 0 \\ 1 \\ 0 \end{pmatrix} = \left(\frac{x_1}{2} + x_2\right) \begin{pmatrix} 2 \\ 0 \\ 1 \\ 0 \end{pmatrix},$$

d.h. die Spalten von A haben die Form $(2c, 0, c, 0)'$ mit $c \in \mathbb{R}$ (Barrantes Campos (2012), p. 231). $\square$

Wenn also eine Matrix $M \in \mathbb{R}^{m,n}$ eine Abbildung definiert, so kann man fragen, ob sie injektiv, surjektiv oder bijektiv ist. Die Abbildung ist injektiv, wenn aus $M\mathbf{x} = \mathbf{u}$ und $X\mathbf{y} = \mathbf{v}$ und $\mathbf{u} = \mathbf{v}$ folgt, dass $\mathbf{x} = \mathbf{y}$ ist, und aus $\mathbf{u} \neq \mathbf{v}$ folgt $\mathbf{x} \neq \mathbf{y}$. Die Frage nach der Injektivität ist also eine Frage nach der Eindeutigkeit der Abbildung. M definiert eine surjektive Abbildung, wenn für jeden Vektor $\mathbf{u} \in V_n$ ein Vektor $\mathbf{x} \in V_m$ existiert derart, dass $M\mathbf{x} = \mathbf{u}$, d.h. die Frage nach der Surjektivität ist die Frage, ob durch M alle Elemente von V_n bestimmt werden. M definiert eine bijektive Abbildung, wenn M eine sowohl injektive wie auch surjektive Abbildung definiert. Dies ist die Frage, ob eine surjektive Abbildung auch eindeutig ist. Offenbar hängen diese Eigenschaften von der Struktur der Matrix M ab. Was mit dem Begriff der Struktur einer Matrix genau gemeint ist, wird im Folgenden entwickelt.

Beispiel 3.3 Es sei

$$T = \begin{pmatrix} t_{11} & t_{12} \\ t_{21} & t_{22} \end{pmatrix} = \begin{pmatrix} \cos \phi & -\sin \phi \\ \sin \phi & \cos \phi \end{pmatrix}. \quad (3.43)$$

T definiert eine Abbildung $\mathbb{R}^2 \rightarrow \mathbb{R}^2$:

$$T\mathbf{x} = \begin{pmatrix} x_1 \cos \phi + x_2 \sin \phi \\ x_1 \sin \phi - x_2 \cos \phi \end{pmatrix} = x_1 \begin{pmatrix} \cos \phi \\ \sin \phi \end{pmatrix} + x_2 \begin{pmatrix} \sin \phi \\ -\cos \phi \end{pmatrix} = \begin{pmatrix} y_1 \\ y_2 \end{pmatrix}.$$

T definiert die Rotation eines Vektors $\mathbf{x}$ um einen Winkel ϕ . Dadurch werden die Elemente von $\mathbb{R}^2$ auf Elemente von $\mathbb{R}^2$ abgebildet, $-\mathbf{y}$ ist ja wieder ein Element von $\mathbb{R}^2$. Die Abbildung ist sicher injektiv und surjektiv, also bijektiv und damit umkehrbar, d.h. man kann einen Vektor $\mathbf{y} \in \mathbb{R}^2$ wählen und in "zurückdrehen", so dass man wieder bei $\mathbf{x}$ landet. Die Abbildung bzw. Matrix, die diese inverse Rotation bewirkt, wird mit T^{-1} bezeichnet. Zu einer Matrix A inverse Matrizen A^{-1} existieren nicht notwendig, worauf in Abschnitt 3.8 zurückgekommen wird.

□

3.4 Anwendung: Mittelwerte und Varianzen

Dieser Abschnitt illustriert die Anwendung des dyadischen Produkts und liefert dabei Darstellungen von Varianz-Kovarianzmatrizen, die gelegentlich bei der Herleitung von Aussagen nützlich sind.

Es sei X eine $(m \times n)$ -Matrix von Messwerten; X_{ij} sei der Messwert beim i -ten Fall für die j -te Variable. Gesucht sind die Mittelwerte und Varianzen für einzelnen Variablen.

Der Mittelwert und die Varianz für die j -te Variable sind durch

$$\bar{x}_j = \frac{1}{m} \sum_{i=1}^m X_{ij}, \quad s_j^2 = \frac{1}{m} \sum_{i=1}^m (X_{ij} - \bar{x}_j)^2 = \frac{1}{m} \sum_{i=1}^m X_{ij}^2 - \bar{x}_j^2 \quad (3.44)$$

gegeben (für kleinere Stichproben sollte bei der Schätzung für s_j^2 natürlich durch $m - 1$ geteilt werden, aber für die hier durchgeführten grundsätzlichen Betrachtungen ist diese Biaskorrektur nicht von Belang). Für die Kovarianz s_{jk} zweier Variablen V_j und V_k hat man

$$s_{jk} = \frac{1}{m} \sum_{i=1}^m (X_{ij} - \bar{x}_j)(X_{ik} - \bar{x}_k) = \frac{1}{m} \sum_{i=1}^m x_{ij}x_{ik} - \bar{x}_j\bar{x}_k, \quad (3.45)$$

mit

$$x_{jk} = X_{jk} - \bar{x}_j, \quad x_{ik} = X_{ik} - \bar{x}_k. \quad (3.46)$$

(3.270) enthält den Ausdruck für s_j^2 als Spezialfall für $j = k$.

Es sei $\vec{1}$ ein m -dimensionaler Vektor, dessen Komponenten alle gleich 1 seien. Man rechnet leicht nach, dass

$$\bar{\mathbf{x}} = \frac{1}{m} X' \vec{1} = \begin{pmatrix} \bar{x}_1 \\ \bar{x}_2 \\ \vdots \\ \bar{x}_n \end{pmatrix} \quad (3.47)$$

der n -dimensionale Vektor der Mittelwerte der n Variablen ist.

Varianz-Kovarianzmatrizen: Die Matrix der s_{jk} läßt sich in Matrixform darstellen. Es sei $\tilde{\mathbf{x}}_i$ der i -te *Zeilenvektor* der Matrix der Abweichungen $X_{ij} - \bar{x}_j$. Das dyadische Produkt von $\tilde{\mathbf{x}}_i$ mit sich selbst ist

$$\tilde{\mathbf{x}}_i \tilde{\mathbf{x}}_i' = (x_{ij}x_{ik}),$$

wobei $(x_{ij}x_{ik})$ die Matrix der Produkte der j -ten Komponente und der k -ten Komponente von $\tilde{\mathbf{x}}_i$ ist, d.h. das Produkt des zentrierten Messwerts der i -ten Person für die j -te Variable mit dem zentrierten Messwert für die k -te Variable. s_{jk} ist die Summe über alle Personen i . Summiert man die dem dyadischen Produkt entsprechenden Matrizen $\tilde{\mathbf{x}}_i \tilde{\mathbf{x}}_i'$ über alle i , erhält man die Matrix aller Kovarianzen s_{jk} :

$$S = \frac{1}{m} \sum_{i=1}^m \tilde{\mathbf{x}}_i \tilde{\mathbf{x}}_i' = \frac{1}{m} \sum_{i=1}^m (X_{ij} - \bar{x}_j)(X_{ik} - \bar{x}_k)' \quad (3.48)$$

bzw.

$$S = \frac{1}{m} X' X - \bar{\mathbf{x}} \bar{\mathbf{x}}' = \frac{1}{m} (X' X - X' \vec{1} \vec{1}' X). \quad (3.49)$$

Dies ist eine allgemeine und oft nützliche Schreibweise für eine Varianz-Kovarianzmatrix, s. aber auch (3.53) weiter unten.

Die Zentrierungsmatrix: Es sei

$$H = I - \frac{1}{m} \vec{1} \vec{1}'. \quad (3.50)$$

H heißt *Zentrierungsmatrix*. H ist offenbar symmetrisch und idempotent, denn

$$H' = (I - \frac{1}{m} \vec{1}\vec{1}')' = I' - (\frac{1}{m} \vec{1}\vec{1}')' = I - \frac{1}{m} \vec{1}\vec{1}' = H, \quad (3.51)$$

und

$$\begin{aligned} HH &= H' = (I - \frac{1}{m} \vec{1}\vec{1}')(I - \frac{1}{m} \vec{1}\vec{1}') \\ &= I - \frac{1}{m} I\vec{1}\vec{1}' - \frac{1}{m} I\vec{1}\vec{1}' + \frac{1}{m^2} \vec{1}\vec{1}'\vec{1}\vec{1}' \\ &= I - \frac{1}{m} I\vec{1}\vec{1}' = H, \end{aligned} \quad (3.52)$$

denn

$$\vec{1}\vec{1}'\vec{1}\vec{1}' = \vec{1}(\vec{1}'\vec{1})\vec{1}' = m\vec{1}\vec{1}'.$$

Die Idempotenz von H erweist sich u.a. als nützlich, wenn Eigenschaften von Varianz-Kovarianzmatrizen betrachtet werden, s. etwa Abschnitt 3.9.5.

Die Varianz-Kovarianzmatrix S läßt sich nun in der Form

$$S = \frac{1}{m} X'HX \quad (3.53)$$

schreiben. Die übliche Biaskorrektur für die Varianz-Kovarianzschätzungen erhält man, indem man diese Gleichung mit $m/(m-1)$ multipliziert:

$$\hat{S} = \frac{m}{m-1} S = \frac{1}{m-1} X'HX. \quad (3.54)$$

Für die Stichprobenkorrelation r_{jk} gilt

$$r_{jk} = \frac{s_{jk}}{s_j s_k}. \quad (3.55)$$

Es sei

$$D = \begin{pmatrix} s_1 & 0 & 0 & 0 \\ 0 & s_2 & 0 & 0 \\ & & \ddots & \\ 0 & 0 & 0 & s_n \end{pmatrix}.$$

Für die Matrix der Korrelationen bzw. der Kovarianzen erhält man dann

$$R = D^{-1}SD^{-1}, \quad S = DRD. \quad (3.56)$$

3.5 Matrizen und Vektorräume

Dieser Abschnitt liefert eine Illustration des allgemeinen Begriffs eines Vektorraums (Abschnitt 8.3) und kann, wenn daran kein Interesse besteht, übersprungen werden.

In Abschnitt 2.4.1 wurde der allgemeine Begriff eines Vektorraums eingeführt. In diesem Abschnitt wird gezeigt, dass die Menge der $(m \times n)$ -Matrizen einen Vektorraum bildet, – eine Matrix ist also ein 'Vektor' in dem Sinne, in dem die Elemente eines Vektorraums als Vektoren bezeichnet werden. Der Definition 8.4 auf Seite 226 entnimmt man, wenn man für $\mathbf{v}$ eine $(m \times n)$ -Matrix A einsetzt und die Regeln für die Multiplikation einer Matrix mit einem Skalar und die Addition von Matrizen berücksichtigt, die Menge der $(m \times n)$ -Matrizen einen Vektorraum über einem Körper $\mathbb{K}$ bilden. Genauer wird diese Aussage in der folgenden Weise ausgedrückt:

Es sei $\mathbb{K}$ ein Körper (hier $\mathbb{K} = \mathbb{R}$ oder $\mathbb{K} = \mathbb{C}$), und $m, n \in \mathbb{N}$, d.h. m und n sind natürliche Zahlen. Dann definiert

$$\mathbb{K}^{m \times n} = \{a_{jk}, j = 1, \dots, m, k = 1, \dots, n | a_{jk} \in \mathbb{K}\} \quad (3.57)$$

die Menge der Menge der $(m \times n)$ -Matrizen. Sind A und B zwei Matrizen aus $\mathbb{K}^{m \times n}$, so ist deren Summe wie in (3.13) durch

$$A + B = (a_{jk} + b_{jk}), \quad j = 1, \dots, m, k = 1, \dots, n$$

erklärt und für $\lambda \in \mathbb{K}$ ist $\lambda A = (\lambda a_{jk})$ erklärt. Damit ist ein Vektorraum $(\mathbb{K}, +, \cdot)$ erklärt.

Das neutrale Element des Vektorraums ist durch $0_{\mathbb{K}}$, also der Null des Körpers $\mathbb{K}$ definiert, und das in Bezug auf '+' inverse Element zu A ist $-A = (-a_{jk})$.

Für jeden Vektorraum existiert eine Basis, und die Frage ist nun, wie im vorliegenden Fall eine Basis definiert werden kann. Dazu kann man die $(m \times n)$ -Matrizen E_{jk} definieren: die Elemente dieser Matrizen sind alle gleich Null bis auf das Element in der j -ten Zeile und k -ten Spalte, das gleich 1 ist:

$$E_{jk} = (e_{uv}), \quad e_{uv} = \begin{cases} 0, & u \neq j \vee v \neq k \\ 1, & u = j \wedge v = k \end{cases} \quad (3.58)$$

wobei $\vee$ für das einschließende 'Oder' steht ($p \vee q$ ist wahr, wenn entweder p oder q wahr ist, oder sowohl p wie q wahr sind), und $\wedge$ steht für die Konjunktion 'und' ($p \wedge q$ ist wahr nur dann, wenn sowohl p wie auch q wahr sind). Man weist leicht nach, dass die E_{jk} linear unabhängig sind (analog zum Nachweis der linearen Unabhängigkeit der Einheitsvektoren $\mathbf{e}_j$). Eine gegebene Matrix $A \in \mathbb{K}^{m \times n}$ ist dann als Linearkombination dieser Matrizen darstellbar:

$$A = \sum_{j=1}^m \sum_{k=1}^n a_{jk} E_{jk}. \quad (3.59)$$

Es gibt $m \cdot n$ Matrizen E_{jk} , so dass die Dimension des Matrizenraums gleich $m \cdot n$ ist.

3.6 Der Rang einer Matrix

Auf Seite 49 ist der Begriff des Ranges eines Vektorraums bzw. eines Teilraums eines Vektorraums eingeführt worden. Es sei nun $X = [\mathbf{x}_1, \dots, \mathbf{x}_n]$ eine $(m \times n)$ -Matrix. Die lineare Hülle $\mathcal{L}_s = \mathcal{L}(\mathbf{x}_1, \dots, \mathbf{x}_n)$ ist ein Vektorraum und habe den Rang s , d.h. die als Linearkombinationen der Spaltenvektoren $\mathbf{x}_j$ erzeugten Vektoren seien als Linearkombinationen von s linear unabhängigen, m -dimensionalen Vektoren $\mathbf{u}_1, \dots, \mathbf{u}_s$ darstellbar, die zu einer Matrix $U = [\mathbf{u}_1, \dots, \mathbf{u}_s]$ zusammengefasst werden können. Die $\mathbf{u}_1, \dots, \mathbf{u}_s$ bilden eine Basis für $\mathcal{L}_s$. Es gibt beliebig viele Basen für $\mathcal{L}_s$, und $\mathbf{u}_1, \dots, \mathbf{u}_s$ kann beliebig gewählt werden.

Die Zeilenvektoren von X sind die n -dimensionalen Spaltenvektoren $\tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_m$ von X' . Die lineare Hülle $\mathcal{L}_z = \mathcal{L}(\tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_m)$ habe den Rang $r < m$, d.h. die $\tilde{\mathbf{x}}_i$ seien als Linearkombinationen von r linear unabhängigen Vektoren $\mathbf{v}_1, \dots, \mathbf{v}_r$ darstellbar, die zu einer Matrix $V = [\mathbf{v}_1, \dots, \mathbf{v}_r]$ zusammengefasst werden können. $\mathbf{v}_1, \dots, \mathbf{v}_r$ ist ebenfalls eine beliebig gewählte Basis für $\mathcal{L}_z$. s heißt *Spaltenrang* von X , und r ist der *Zeilenrang* von X . Man beachte, dass die Spaltenvektoren von X m -dimensional, die Zeilenvektoren von X aber n -dimensional sind, d.h. die Zeilen- und Spaltenvektoren sind Elemente aus Vektorräumen mit verschiedener Dimensionalität.

Die Darstellungen der Spaltenvektoren $\mathbf{x}_j$ von X und der Spaltenvektoren $\tilde{\mathbf{x}}_i$ von X' sind nicht unabhängig von einander. So sei $X = UA'$, $U = [\mathbf{u}_1, \dots, \mathbf{u}_s]$, und A' die Matrix der Koeffizienten: der Spaltenvektor $\mathbf{a}_j$ von A' ist der Koeffizientenvektor für $\mathbf{x}_j$: $\mathbf{x}_j = U\mathbf{a}_j$. Dann folgt aber $X' = AU'$, d.h. die Spaltenvektoren $\tilde{\mathbf{x}}_i$ sind Linearkombinationen der Spaltenvektoren von A , und U' ist nun die Koeffizientenmatrix. A ist eine $(n \times s)$ -Matrix. Dem obigen Ansatz zufolge soll aber $X' = VB'$ gelten, wobei B' die zugehörige Matrix der Koeffizienten für die $\mathbf{v}_1, \dots, \mathbf{v}_r$ ist. Man könnte nun einfach folgern, dass die oben angenommene Matrix V gleich der Matrix A ist und $B = U'$ gilt, und dass somit $r = s$ gilt, aber diese Gleichsetzung setzt voraus, dass die Spaltenvektoren von A linear unabhängig sind. Diese Eigenschaft muß allerdings noch nachgewiesen werden. Dieser Nachweis wird im Beweis des folgenden Satzes geliefert.

Satz 3.2 *Es sei X eine $(m \times n)$ -Matrix mit dem Zeilenrang r und dem Spaltenrang s . Dann gilt*

$$r = s \quad (3.60)$$

d.h. der Zeilenrang ist stets gleich dem Spaltenrang, sowie

$$X = UV', \quad (3.61)$$

wobei $U \in \mathbb{R}^{m,r}$ und $V \in \mathbb{R}^{n,r}$ Matrizen mit dem Rang r sind. Weiter gilt

$$r \leq \min(m, n) \quad (3.62)$$

Beweis: X besteht aus den n m -dimensionalen Spaltenvektoren $\mathbf{x}_1, \dots, \mathbf{x}_n$, bzw. aus den m n -dimensionalen Zeilenvektoren $\tilde{\mathbf{x}}'_1, \dots, \tilde{\mathbf{x}}'_m$:

$$X = [\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n] = \begin{bmatrix} \tilde{\mathbf{x}}'_1 \\ \tilde{\mathbf{x}}'_2 \\ \vdots \\ \tilde{\mathbf{x}}'_m \end{bmatrix}.$$

Die Spaltenvektoren $\mathbf{x}_j$ können als Linearkombinationen von s linear unabhängigen m -dimensionalen Vektoren $\mathbf{u}_1, \dots, \mathbf{u}_s$ angeschrieben werden. Fasst man die $\mathbf{u}_k$, $k = 1, \dots, s$ zu einer Matrix $U = [\mathbf{u}_1, \dots, \mathbf{u}_s]$ zusammen, so erhält man

$$X = [\mathbf{x}_1, \dots, \mathbf{x}_n] = UA', \quad (3.63)$$

wobei A' die zu U gehörende $(s \times n)$ -Matrix der Koeffizienten ist (s. Satz 2.1 über die Eindeutigkeit der Koeffizienten bei linear unabhängigen $\mathbf{u}_j$, Seite 32): die Spaltenvektoren $\mathbf{a}_j = (a_{1j}, \dots, a_{sj})'$ von A' erzeugen $\mathbf{x}_j$ gemäß

$$\mathbf{x}_j = a_{1j}\mathbf{u}_1 + \dots + a_{sj}\mathbf{u}_s = U\mathbf{a}_j$$

Die Spaltenvektoren $\tilde{\mathbf{x}}_i$ von X' (also die Zeilenvektoren von X) können als Linearkombinationen von r linear unabhängigen Spaltenvektoren $\mathbf{v}_1, \dots, \mathbf{v}_r$ dargestellt werden, wobei die $\mathbf{v}_k$, $1 \leq k \leq r$ zunächst beliebig gewählt werden können unter der Einschränkung, dass sie linear unabhängig sind. Fasst man sie zu einer Matrix V zusammen, so erhält man

$$X' = [\tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_m] = VB', \quad (3.64)$$

und B' ist die zu V korrespondierende Koeffizientenmatrix, analog zu A' in (3.63), und

$$\tilde{x}_i = b_{1i}\mathbf{v}_1 + \dots + b_{ri}\mathbf{v}_r = V\mathbf{b}_i.$$

Da U in (3.63) s Spalten hat, muß A' s Zeilen haben. Die Zeilenvektoren von X können aber auch als Linearkombinationen der Zeilenvektoren von A' aufgefasst werden (vergl. Punkt 2, Seite 68). Dann folgt, dass der Zeilenrang r von X nicht größer als s sein kann (es gibt eben nur s Zeilen in A'), d.h. es muß $r \leq s$ gelten.

Umgekehrt folgt aus (3.64), dass die Zeilenvektoren von X' (also die Spaltenvektoren von X) auch als Linearkombinationen der Zeilenvektoren von B' dargestellt werden können. Da V nach Voraussetzung r Spalten hat, muß B' r Zeilen haben. Daraus folgt, dass der Zeilenrang s (d.h. der Spaltenrang von X) von X' nicht größer als r sein kann (es gibt eben nur r Zeilen in B'), d.h. es muß $s \leq r$ gelten. Beide Ergebnisse zusammen liefern

$$r \leq s \wedge s \leq r \Rightarrow r = s,$$

wie behauptet¹⁵.

Nach (3.63) gilt $X = UA'$, und nach (3.64) gilt $X' = VB'$, also $X = BV'$. U und B haben notwendig denselben Rang r ; B repräsentiert möglicherweise eine von U verschiedene Basis von $\mathcal{L}_s$. Da U und B beliebig gewählt werden können, kann man insbesondere $B = U$ wählen. Man hat dann $X = UA' = UV'$, und da die Linearkombinationen für die $\mathbf{x}_j$ eindeutig sind (Satz 2.1, Seite 32), wenn die Spaltenvektoren von U linear unabhängig sind, hat man mit der Wahl von U auch die von A festgelegt, denn es muß nun $A' = V'$ gelten, womit nachgewiesen ist, dass X als Produkt UV' zweier Matrizen mit gleichem Rang dargestellt werden kann.

U und V haben wegen $r = s$ jeweils r linear unabhängige Spaltenvektoren. Es sei $m > n$. Da der Rang von X gleich der Maximalzahl der linear unabhängigen Vektoren von $\mathcal{L}(\mathbf{x}_1, \dots, \mathbf{x}_n)$ und von $\mathcal{L}(\tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_m)$ ist, folgt $r \leq n$. Analog folgt für $m < n$, dass $r \leq m$ sein muß. Zusammengefaßt ergibt sich die Aussage

$$r \leq \min(m, n).$$

also (3.66). □

Anmerkungen:

1. U ist eine $(m \times r)$ -Matrix und V ist eine $(n \times r)$ -Matrix; beide Matrizen haben r Spalten, die die latenten Variablen repräsentieren.
2. Für $r = \min(m, n)$ sagt man, X habe den *vollen Rang*.
3. Aus dem Beweis zu Satz 3.2 ergibt sich, dass sich *jede* $(m \times n)$ -Matrix X als Produkt der Form (3.61) darstellen läßt.
4. Die in Abschnitt 3.9.8 eingeführte Singularwertzerlegung von X ist ein Spezialfall von (3.61), der vielen faktorenanalytischen Ansätzen zur Interpretation von Datenmatrizen zugrunde liegt. □

Folgerungen:

1. Aus (3.61) folgen sofort die Aussagen

$$\text{rg}(X'X) = \text{rg}(XX') = \text{rg}(X), \quad (3.65)$$

denn $X'X = V'U'UV = V'C$, $C = U'UV$, d.h. die Spalten von $X'X$ sind Linearkombinationen der r linear unabhängigen n -dimensionalen Spaltenvektoren von V' , und $XX' = UVV'U = UD$, $D = VV'U$, und die Spalten von XX' sind Linearkombinationen der r linear unabhängigen m -dimensionalen Spaltenvektoren von U . In Abschnitt 3.9.8, Seite ?? (Satz ??) wird ein weiterer Beweis dieser Aussage gegeben. Die Aussagen (3.65) sind für die multivariate Statistik von Bedeutung: bei geeigneter Zentrierung entsprechen den Matrizen $X'X$ bzw. XX' Kovarianz- bzw. Korrelationsmatrizen, die demnach denselben Rang wie die Datenmatrix haben.

¹⁵Das Zeichen $\wedge$ steht für 'und'.

2. Es sei $m > n$. Da der Rang von X gleich der Maximalzahl der linear unabhängigen Vektoren von $\mathcal{L}(\mathbf{x}_1, \dots, \mathbf{x}_n)$ und von $\mathcal{L}(\tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_m)$ ist, folgt $r \leq n$. Analog folgt für $m < n$, dass $r \leq m$ sein muß. Zusammengefaßt ergibt sich die Aussage

$$r \leq \min(m, n). \quad (3.66)$$

Für $r = \min(m, n)$ sagt man, X habe den *vollen Rang*.

Da stets der Zeilen- gleich dem Spaltenrang einer Matrix ist, genügt es, nur vom Rang einer Matrix zu sprechen:

Definition 3.5 *Es sei r die Maximalzahl der linear unabhängigen Spalten- bzw. Zeilenvektoren einer Matrix X . Dann heißt r der Rang der Matrix X .*

Anmerkung: Ein alternativer Beweis macht von den in Abschnitt 2.4.2, Seite 56, eingeführten elementaren Umformungen Gebrauch. Sie lassen sich auf die Matrix X anwenden und transformieren X in eine Matrix E , deren Zeilen und Spalten nur aus Nullen besteht mit Ausnahme der ersten r Elemente e_{ii} , $i = 1, \dots, r$; da die elementaren Umordnungen den Rang von X nicht verändern, ist der Rang von X gleich dem Rang von E , und E enthält r linear unabhängige Zeilen- und r linear unabhängige Spaltenvektoren, d.h. der Zeilen- und Spaltenrang von X muß identisch sein.

Es sei X eine $(m \times n)$ -Matrix mit dem unbekannten Rang $\text{rg}(X) = r \leq \min(m, n)$. Die Frage ist nun, wie der Wert von r bestimmt werden kann. Dazu kann man Umformungen der Zeilen bzw. Spalten von X Gebrauch machen, die einerseits den Rang der Matrix nicht verändern, die die Matrix X in eine Form X_T bringen, von der der Rang sofort abgelesen werden kann. Die letzten $m - r$ Zeilen von X_T enthalten nur Nullen, und in den ersten r Zeilen enthalten alle Zellen bis zur Zahl $\bullet$ (das sogenannte Pivot-Element) nur Nullen, und anschließend durch $*$ gekennzeichnete Zahlen ungleich Null:

$$X_T = \begin{pmatrix} 0 & \dots & 0 & \bullet & * & * & * & * & * & * & \dots & * & \dots & * \\ 0 & \dots & 0 & 0 & 0 & \bullet & * & * & * & * & \dots & * & \dots & * \\ 0 & \dots & 0 & 0 & 0 & 0 & \bullet & * & * & * & \dots & * & \dots & * \\ 0 & \dots & 0 & 0 & 0 & 0 & 0 & 0 & \bullet & * & \dots & * & \dots & * \\ \vdots & & \vdots & \vdots & \vdots & \vdots & \vdots & & \ddots & \vdots & & \vdots & & \vdots \\ 0 & \dots & 0 & 0 & 0 & 0 & 0 & 0 & 0 & \bullet & \dots & * & \dots & * \\ 0 & \dots & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & \dots & 0 & \dots & 0 \\ \vdots & & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & & \vdots \\ 0 & \dots & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & \dots & 0 & \dots & 0 \end{pmatrix} \quad (3.67)$$

Aus dem Beweis von Satz 3.2 folgt, dass eine beliebige $(m \times n)$ -Matrix X vom Rang r stets in der Form

$$X = UV \quad (3.68)$$

repräsentiert werden kann, wobei U eine $(m \times r)$ -Matrix und V eine $(r \times n)$ -Matrix ist, und U und V haben den Rang r . Denn wenn die Spaltenvektoren von X Linearkombinationen der Spaltenvektoren von U , so enthält V die notwendigen Koeffizienten. Die Zeilenvektoren von X sind dann aber Linearkombinationen der Zeilenvektoren von V , und der Rang von V muß ebenfalls gleich r sein: der Rang von V kann nicht größer als r sein, allenfalls kleiner, aber da der Zeilenrang von X ebenfalls gleich r ist, folgt, dass V den Rang r , denn wäre der Rang kleiner als r , so wäre der Zeilenrang von X kleiner als r , was nicht möglich. In analoger Weise folgt, dass der Rang von U gleich r sein muß, wenn V den Rang r hat. Die Repräsentation (3.68) erweist sich als grundlegend für viele Verfahren der multivariaten Statistik.

Satz 3.3 *Es sei A eine $(m \times n)$ -Matrix, B sei eine $(n \times p)$ -Matrix. Dann ist das Produkt $C = AB$ eine $(m \times p)$ -Matrix. Der Rang von C ist kleiner oder gleich dem kleineren der Ränge von A und B , d.h.*

$$\text{rg}(C) \leq \min[\text{rg}(A), \text{rg}(B)]. \quad (3.69)$$

Beweis: A habe den Rang s . Die Spaltenvektoren von C sind Linearkombinationen der Spaltenvektoren von A , d.h. $C \subseteq \mathcal{L}(A) = \mathcal{L}(\mathbf{a}_1, \dots, \mathbf{a}_s)$, $\mathbf{a}_1, \dots, \mathbf{a}_s$ eine Basis der linearen Hülle $\mathcal{L}(A)$ von A , so dass C höchstens den Spaltenrang s hat. Die Matrix B habe den Rang r , und die Zeilenvektoren von C sind Linearkombinationen der Zeilenvektoren von B , so dass $C' \subseteq \mathcal{L}(B') = \mathcal{L}(\mathbf{b}_1, \dots, \mathbf{b}_r)$, $\mathbf{b}_1, \dots, \mathbf{b}_r$ eine Basis von $\mathcal{L}(B')$, so dass C höchstens den Rang r haben kann. Es sei $r \leq s$; dann hat C höchstens den Rang r , denn nach Satz 3.60 sind dann auch die Spaltenvektoren von C als Linearkombinationen von maximal r linear unabhängigen Vektoren darstellbar. Sei umgekehrt $s \leq r$; analog zur vorangehenden Argumentation ist dann ist der Rang von C höchstens gleich s , d.h. $\text{rg}(C) \leq \min(\text{rg}(A), \text{rg}(B))$. $\square$

Satz 3.4 *Es sei $\mathbf{x}$ ein m -dimensionaler Vektor und $\mathbf{y}$ sei ein n -dimensionaler Vektor. Dann hat das dyadische Produkt $\mathbf{x}\mathbf{y}'$ den Rang 1.*

Beweis: Es genügt, die Spalten der Matrix $\mathbf{x}\mathbf{y}'$ zu betrachten. Die j -te und die k -te Spalte sind durch

$$\begin{pmatrix} x_1 y_j \\ x_2 y_j \\ \vdots \\ x_m y_j \end{pmatrix} = y_j \mathbf{x}, \quad \begin{pmatrix} x_1 y_k \\ y_2 y_k \\ \vdots \\ x_m y_k \end{pmatrix} = y_k \mathbf{x}$$

gegeben. Die Spaltenvektoren sind also alle parallel zueinander. Mithin hat $\mathbf{x}\mathbf{y}'$ den Rang 1. $\square$

3.7 Determinanten

3.7.1 Die Definition der Determinante

Gegeben sei ein lineares Gleichungssystem mit zwei Unbekannten x_1 und x_2 :

$$a_{11}x_1 + a_{12}x_2 = y_1 \quad (3.70)$$

$$a_{21}x_1 + a_{22}x_2 = y_2 \quad (3.71)$$

Löst man (3.70) nach x_2 auf und substituiert den resultierenden Ausdruck in (3.71), so findet man einen Ausdruck für x_1 ; setzt man diesen in (3.70) ein, so erhält man eine Lösung für x_1 . Man findet

$$x_1 = \frac{a_{22}y_1 - a_{12}y_2}{a_{11}a_{22} - a_{12}a_{21}} \quad (3.72)$$

$$x_2 = \frac{a_{11}y_2 - a_{21}y_1}{a_{11}a_{22} - a_{12}a_{21}} \quad (3.73)$$

Der Nenner hat in beiden Ausdrücke dieselbe Form; er ist die *Determinante* der Koeffizientenmatrix A . Man schreibt

$$|A| = \det(A) = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} = a_{11}a_{22} - a_{12}a_{21}. \quad (3.74)$$

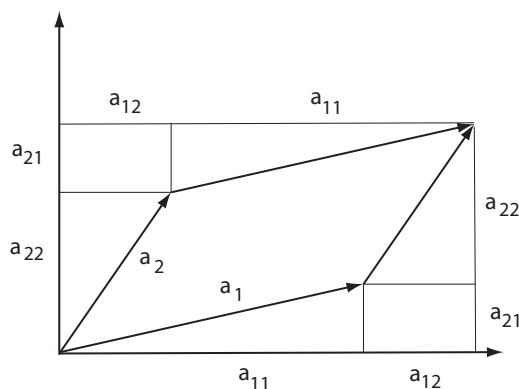
Die ersten drei Ausdrücke $|A|$ und $\det(A)$ sind äquivalente Bezeichnungen für die Determinante von A . Wie die rechte Seite zeigt, ist $|A|$ durch die Differenz der Produkte der Elemente in der Hauptdiagonalen (a_{11}, a_{22}) und dem Produkt der Elemente in der Nebendiagonalen (a_{12}, a_{21}) gegeben. Die Determinante ist zunächst nur eine Schreibweise für den Nenner der Lösungen für ein Gleichungssystem mit zwei Unbekannten.

Offenbar ist $|A|$ gleich dem Flächeninhalt des durch die Vektoren $\mathbf{a}_1$ und $\mathbf{a}_2$ definierten Parallelogramms, vergl. Abbildung 6. Der Flächeninhalt ist gleich dem Inhalt des Rechtecks mit den Seitenlängen $a_{11} + a_{12}$ und $a_{21} + a_{22}$, minus den Inhalten der in Abb. 6 erklärten Dreiecke und Rechtecke, wie man nur Nachrechnen bestätigt. Man vergleiche diese Abbildung mit Abb. 15 auf Seite 219 und setze $\mathbf{a}_1 = \mathbf{x}$, $\mathbf{a}_2 = \mathbf{y}$; ein alternativer Ausdruck für den Flächeninhalt ist dann $F = \|\mathbf{a}_1\| \|\mathbf{a}_2\| \sin \phi$, wobei ϕ der von $\mathbf{a}_1$ und $\mathbf{a}_2$ definierte Winkel ist. Die Beziehung zwischen Determinante und Volumen verallgemeinert sich auf $(n \times n)$ -Matrizen: Die Determinante ist gleich dem Volumen des durch die Spaltenvektoren aufgespannten Parallelepipeds.

Es liegt nahe, auch die Zähler in (3.72) und (3.73) als Determinanten auszudrücken, wenn man nur die zugehörigen Matrizen geeignet definiert. Definiert man

$$A_1 = \begin{pmatrix} y_1 & a_{12} \\ y_2 & a_{22} \end{pmatrix}, \quad (3.75)$$

Abbildung 6: Zum Flächeninhalt (Determinante) $F = a_{11}a_{22} - a_{12}a_{21}$ des durch die Vektoren $\mathbf{a}_1$ und $\mathbf{a}_2$ definierten Parallelogramms



und berechnet die Determinante $|A_1|$ wie die von $|A|$, so erhält man

$$|A_1| = a_{22}y_1 - a_{12}y_2. \quad (3.76)$$

Dies ist der Zähler in der Lösung für x_1 . Für

$$A_2 = \begin{pmatrix} a_{11} & y_1 \\ a_{21} & y_2 \end{pmatrix} \quad (3.77)$$

findet man

$$|A_2| = a_{11}y_2 - a_{21}y_1, \quad (3.78)$$

und dies ist der Zähler in der Lösung für x_2 . Man kann die Lösungen (3.72) und (3.73) dann in der Form

$$x_1 = \frac{|A_1|}{|A|}, \quad x_2 = \frac{|A_2|}{|A|} \quad (3.79)$$

schreiben. Der Vorzug dieser Schreibweise ist, dass sie die Systematik, nach der die Lösungen gefunden werden kann, angibt, insbesondere, wenn sie sich allgemein für Gleichungssysteme mit n Unbekannten ergibt, d.h. wenn man allgemein Determinanten für $(n \times n)$ -Matrizen definieren kann. Das kann man, und die Lösungen (3.79) sind ein Beispiel für die *Cramersche Regel*¹⁶. Die Matrix A_1 ergab sich, indem der erste Spaltenvektor von A durch den Vektor $\mathbf{y}$ ersetzt wurde, und A_2 ergab sich, indem der zweite Spaltenvektor durch $\mathbf{y}$ ersetzt wurde. Diese Regel gilt allgemein: in einem System mit n Unbekannten findet man für die Unbekannte x_j

$$x_j = \frac{|A_j|}{|A|},$$

¹⁶Nach Gabriel Cramer (1704 – 1752), Genfer Mathematiker, der diese Regel zum ersten Mal herleitete.

wobei aber die Matrizen A_j und A ($n \times n$)-Matrizen sind und A_j aus A entsteht, indem man den j -ten Spaltenvektor von A durch den (nun n -dimensionalen) Vektor $\mathbf{y}$ ersetzt.

Um die Cramersche Regel konkret anwenden zu können, muß ein Ausdruck für die Determinante einer allgemeinen ($n \times n$)-Matrix gefunden werden. Eine Möglichkeit ist, ein Gleichungssystem mit drei Unbekannten zu lösen; im Nenner der Lösungen für x_1, x_2 und x_3 sollte dann ein Ausdruck auftreten, der für die drei Lösungen identisch ist und sich aus den Elementen der Koeffizientenmatrix A zusammensetzt. Das soll hier nicht im Detail durchgeführt werden, aber die Lösung für $|A|$ (und damit implizit auch die für $|A_1|, |A_2|$ und $|A_3|$) kann angegeben werden:

$$\begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} = -a_{21} \begin{vmatrix} a_{12} & a_{13} \\ a_{32} & a_{33} \end{vmatrix} + a_{22} \begin{vmatrix} a_{11} & a_{13} \\ a_{31} & a_{33} \end{vmatrix} - a_{23} \begin{vmatrix} a_{11} & a_{12} \\ a_{31} & a_{32} \end{vmatrix} \quad (3.80)$$

(vergl. Satz 3.7). Die Regel, nach der hier $|A|$ berechnet wurde, ist leicht zu sehen. Die zweite Zeile von A enthält die Elemente a_{21}, a_{22} und a_{23} . Im ersten Summanden auf der rechten Seite taucht a_{21} als Faktor der Determinante einer (2×2)-Matrix auf; diese Matrix entsteht, indem man in der Matrix A die zweite Zeile und die erste Spalte streicht. Der zweite Summand hat den Faktor a_{22} der Determinante einer Matrix, die aus A entsteht, wenn man wieder die zweite Zeile, jetzt aber die zweite Spalte streicht. Der dritte Faktor ist a_{23} , und die Determinante ist die der Matrix, die entsteht, wenn man wiederum die zweite Zeile, jetzt aber die dritte Spalte aus A streicht. Wie man die Determinante einer (2×2)-Matrix berechnet, ist ja bereits bekannt. Es bleibt noch zu klären, wie die Vorzeichen zu finden sind. Der Vorzeichen für den Faktor a_{ij} ist $(-1)^{i+j}$, für a_{21} findet man also $(-1)^3 = -1$, für a_{22} findet man $(-1)^4 = +1$, und für a_{23} erhält man $(-1)^5 = -1$. Man sagt, die Determinante ist nach der zweiten Zeile von A entwickelt worden. Das Ergebnis ist numerisch identisch, wenn man sie nach der ersten oder der dritten Zeile entwickelt. Die Determinanten in (3.80) heißen *Kofaktoren* der Faktoren a_{21}, a_{22} und a_{23} . Berechnet man die Determinante für eine ($n \times n$)-Matrix, so geht man analog vor: man wählt eine Zeile, bestimmt die Vorzeichen nach Maßgabe der Indices in der Zeile, und berechnet die Kofaktoren als Determinanten für $(n-1 \times n-1)$ -Matrizen, und diese wiederum werden auf Determinanten von $(n_2 \times n_2)$ -Matrizen zurückgeführt. Dass Determinanten tatsächlich auf diese Weise berechnet werden können, kann man herausfinden, indem man entsprechende Gleichungssysteme analysiert. Es sei noch angemerkt, dass sich Determinanten auch als Maße für das Volumen des Parallelepiped ergeben, dass durch die Spaltenvektoren von A aufgespannt wird; auf diese Weise gelangt auch die Determinante $|\Sigma|$ der Varianz-Kovarianzmatrix einer n -dimensionalen Normalverteilung in den Normierungsfaktor dieser Verteilung.

Der allgemeine Begriff der Determinante wird nach Weierstraß¹⁷ axiomatisch

¹⁷Karl Weierstraß (1815 – 1897), Mathematiker

eingeführt. Dazu wird zunächst noch einmal festgestellt, dass die Determinante $\det(A)$ einer Matrix A eine reelle Zahl ist. Weiter seien $\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n$ die Zeilenvektoren von A , so dass

$$A = \begin{pmatrix} \mathbf{a}_1 \\ \mathbf{a}_2 \\ \vdots \\ \mathbf{a}_n \end{pmatrix}.$$

Dann

Definition 3.6 Die Determinante einer $n \times n$ -Matrix A mit den Zeilenvektoren $\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n$ ist eine Abbildung $A \rightarrow \det(A) \in \mathbb{R}$ mit den Eigenschaften

1. $\det(A)$ ist linear in jeder Zeile,

(a) d.h. gilt für ein i , $1 \leq i \leq n$ $\mathbf{a}_i = \mathbf{a}'_i + \mathbf{a}''_i$, so ist

$$\det \begin{pmatrix} \vdots \\ \mathbf{a}_i \\ \vdots \end{pmatrix} = \det \begin{pmatrix} \vdots \\ \mathbf{a}'_i \\ \vdots \end{pmatrix} + \det \begin{pmatrix} \vdots \\ \mathbf{a}''_i \\ \vdots \end{pmatrix}. \quad (3.81)$$

(b) Ist $\mathbf{a}_i = \lambda \mathbf{a}'_i$ so ist

$$\det \begin{pmatrix} \vdots \\ \mathbf{a}_i \\ \vdots \end{pmatrix} = \lambda \det \begin{pmatrix} \vdots \\ \mathbf{a}'_i \\ \vdots \end{pmatrix} \quad (3.82)$$

An den mit Punkten bezeichneten Stellen stehen die Vektoren

$$\mathbf{a}_1, \dots, \mathbf{a}_{i-1}, \mathbf{a}_{i+1}, \dots, \mathbf{a}_n.$$

2. $\det$ ist alternierend, dh hat A zwei gleiche Zeilen, so ist $\det(A) = 0$.

3. $\det$ ist normiert, dh $\det(I_n) = 1$, I_n die $n \times n$ -Einheitsmatrix.

3.7.2 Eigenschaften der Determinante

Satz 3.5 Es sei A eine $n \times n$ -Matrix. Dann gelten die folgenden Aussagen:

A1 Es sei $\lambda \in \mathbb{R}$; dann ist $\det(\lambda A) = \lambda^n \det(A)$.

A2 Für (mindestens) ein i gelte $\mathbf{a}_i = \vec{0}$. Dann folgt $\det(A) = 0$.

A3 Die Matrix B entstehe aus der Matrix A durch Vertauschung der Zeilen $\mathbf{a}_i$ und $\mathbf{a}_j$. Dann folgt $\det(B) = -\det(A)$.

A4 Die Matrix B entstehe aus der Matrix A , indem der Zeilenvektor $\mathbf{a}_i$ durch $\mathbf{a}_i + \lambda \mathbf{a}_j$ ersetzt wird, wobei $i \neq j$. Dann ist $\det(B) = \det(A)$.

A5 A sei eine obere Dreiecksmatrix, d.h. es sei

$$A = \begin{pmatrix} \lambda_1 & \cdots & a_{1n} \\ 0 & \lambda_2 & \vdots \\ \vdots & \cdots & \vdots \\ 0 & \cdots & \lambda_n \end{pmatrix}. \quad (3.83)$$

Dann ist $\det(A) = \lambda_1 \lambda_2 \cdots \lambda_n$.

A6 $\det(A) = 0$ genau dann, wenn der Rang r von A kleiner als n ist. A und B seien beide $n \times n$ -Matrizen. Dann gilt $\det(AB) = \det(A) \det(B)$.

A7 Es ist $\det(A') = \det(A)$ und $\det(A^{-1}) = \det(A)^{-1}$, wenn $r = n$.

Beweis: [A1] folgt sofort aus (3.82), denn λA bedeutet ja, dass jede Zeile (und damit jedes Element) von A mit λ multipliziert wird. [A2] folgt wiederum aus (3.82), denn es sei $\mathbf{a}_i = 0$; dann kann man dafür $\lambda \mathbf{a}_i = \lambda 0$ mit $\lambda = 0$ schreiben. Um [A3] einzusehen, erinnere man sich an Punkt 2: hat eine Matrix zwei identische Zeilen, so ist die Determinante der Matrix gleich Null, d.h.

$$\begin{aligned} 0 &= \det \begin{pmatrix} \mathbf{a}_1 \\ \vdots \\ \mathbf{a}_i + \mathbf{a}_j \\ \vdots \\ \mathbf{a}_i + \mathbf{a}_j \\ \vdots \\ \mathbf{a}_n \end{pmatrix} \stackrel{1}{=} \det \begin{pmatrix} \mathbf{a}_1 \\ \vdots \\ \mathbf{a}_i \\ \vdots \\ \mathbf{a}_i + \mathbf{a}_j \\ \vdots \\ \mathbf{a}_n \end{pmatrix} + \det \begin{pmatrix} \mathbf{a}_1 \\ \vdots \\ \mathbf{a}_j \\ \vdots \\ \mathbf{a}_i + \mathbf{a}_j \\ \vdots \\ \mathbf{a}_n \end{pmatrix} \\ &\stackrel{1}{=} \det \begin{pmatrix} \mathbf{a}_1 \\ \vdots \\ \mathbf{a}_i \\ \vdots \\ \mathbf{a}_i \\ \vdots \\ \mathbf{a}_n \end{pmatrix} + \det \begin{pmatrix} \mathbf{a}_1 \\ \vdots \\ \mathbf{a}_i \\ \vdots \\ \mathbf{a}_j \\ \vdots \\ \mathbf{a}_n \end{pmatrix} + \det \begin{pmatrix} \mathbf{a}_1 \\ \vdots \\ \mathbf{a}_j \\ \vdots \\ \mathbf{a}_i \\ \vdots \\ \mathbf{a}_n \end{pmatrix} + \det \begin{pmatrix} \mathbf{a}_1 \\ \vdots \\ \mathbf{a}_j \\ \vdots \\ \mathbf{a}_j \\ \vdots \\ \mathbf{a}_n \end{pmatrix} \end{aligned}$$

$$\stackrel{1}{=} \det \begin{pmatrix} \mathbf{a}_1 \\ \vdots \\ \mathbf{a}_i \\ \vdots \\ \mathbf{a}_j \\ \vdots \\ \mathbf{a}_n \end{pmatrix} + \det \begin{pmatrix} \mathbf{a}_1 \\ \vdots \\ \mathbf{a}_j \\ \vdots \\ \mathbf{a}_i \\ \vdots \\ \mathbf{a}_n \end{pmatrix} = 0,$$

woraus [A3] folgt (die Zahlen über den Gleichheitszeichen geben die Annahmen bzw Axiome an). Für [A5] schließlich hat man

$$\det \begin{pmatrix} \mathbf{a}_1 \\ \vdots \\ ba_i + ba_j \\ \vdots \\ \mathbf{a}_j \\ \vdots \\ \mathbf{a}_n \end{pmatrix} \stackrel{1}{=} \det \begin{pmatrix} \mathbf{a}_1 \\ \vdots \\ ba_i \\ \vdots \\ \mathbf{a}_j \\ \vdots \\ \mathbf{a}_n \end{pmatrix} + \lambda \det \begin{pmatrix} \mathbf{a}_1 \\ \vdots \\ \mathbf{a}_j \\ \vdots \\ \mathbf{a}_j \\ \vdots \\ \mathbf{a}_n \end{pmatrix} \stackrel{2}{=} \det \begin{pmatrix} \mathbf{a}_1 \\ \vdots \\ ba_i \\ \vdots \\ \mathbf{a}_j \\ \vdots \\ \mathbf{a}_n \end{pmatrix} + \lambda 0 = 0$$

□

Satz 3.6 Sind die Zeilen von A linear abhängig, so ist $\det(A) = 0$, d.h. die Determinante einer Matrix ist nur dann ungleich Null, wenn die Matrix vollen Rang hat.

Beweis: Sind die Zeilenvektoren von A linear abhängig, so gilt $\sum_j \lambda_j \mathbf{a}_j = 0$ und nicht alle λ_j sind gleich Null. Es sei $\lambda_1 \neq 0$, was man durch Umordnen der Zeilen stets erreichen kann. Dann gilt

$$\mathbf{a}_1 = \sum_{j=2}^n \frac{\lambda_j}{\lambda_1} \mathbf{a}_j,$$

und

$$|A| = \det(A) = \det \begin{pmatrix} -\sum_{j=2}^n (\lambda_j/\lambda_1) \mathbf{a}_j \\ \vdots \\ \mathbf{a}_n \end{pmatrix} \stackrel{(1)}{=} - \sum_{j=2}^n \frac{\lambda_j}{\lambda_1} \det \begin{pmatrix} \mathbf{a}_i \\ \vdots \\ \mathbf{a}_n \end{pmatrix} \stackrel{2}{=} 0.$$

□

Definition 3.7 Es sei A_{ij} die Matrix, die durch Streichen der i -ten Zeile und der j -ten Spalte aus A entsteht. Dann heißt $\tilde{a}_{ij} = (-1)^{i+j} |A_{ij}|$ der (i, j) -te Kofaktor von A . Die Matrix $A^* = (\tilde{a}_{ij})$ heißt die zu A adjungierte Matrix.

Es gilt der

Satz 3.7 (Laplacescher Entwicklungssatz) *Es gelten die Beziehungen*

$$|A| = \sum_{i=1}^n a_{ij} \tilde{a}_{ij} \quad (3.84)$$

$$= \sum_{j=1}^n a_{ij} \tilde{a}_{ij} \quad (3.85)$$

wobei (3.84) die Entwicklung nach der j -ten Spalte und (3.85) die Entwicklung nach der i -ten Zeile darstellt.

Beweis: Illustration für den Fall $n = 3$ s. (3.80). □

Als direkte Folge des Laplaceschen Entwicklungssatzes 3.7 gilt

$$A(A^*)' = |A|I = \begin{pmatrix} |A| & 0 & \cdots & 0 \\ 0 & |A| & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & |A| \end{pmatrix}. \quad (3.86)$$

Diese Beziehung führt auf den Begriff der Inversen einer Matrix (vergl. (3.90) im folgenden Abschnitt).

3.8 Die Inverse einer $(n \times n)$ -Matrix

Für $0 \neq x \in \mathbb{R}$ existiert eine zu x inverse Zahl x^{-1} derart, dass $x \cdot x^{-1} = x^{-1} \cdot x = 1$. Natürlich ist $x^{-1} = 1/x$, weshalb auch $x \neq 0$ vorausgesetzt werden muß: für $x = 0$ ist $1/x$ ja nicht definiert.

Die Einführung einer zu einer gegebenen Matrix A inversen Matrix A^{-1} erweist sich ebenfalls als nützlich, wobei aber nicht einfach $A^{-1} = 1/A$ gesetzt werden kann, – ein solcher Ausdruck macht keinen Sinn. Es zeigt sich auch, dass nicht für *jede* Matrix eine inverse Matrix, die analog zur inversen reellen Zahl x^{-1} definiert ist ($x^{-1} \cdot x = x \cdot x^{-1} = 1$) definierbar ist:

Definition 3.8 *Es seien A und B $(n \times n)$ -Matrizen derart, dass*

$$BA = AB = I. \quad (3.87)$$

Dann heißt B die zu A inverse Matrix; man schreibt

$$B = A^{-1}. \quad (3.88)$$

Es gilt der

Satz 3.8 Es sei $B = A^{-1}$; dann ist A quadratisch. Die Inverse A^{-1} existiert für die eine $(n \times n)$ -Matrix A genau dann, wenn A den Rang n hat, d.h. wenn die Zeilen- bzw. Spaltenvektoren von A linear unabhängig sind.

Beweis: Es sei A eine $(m \times n)$ -Matrix und es existiere A^{-1} . Dann muß A^{-1} eine $(n \times m)$ -Matrix sein, damit $A^{-1}A = I$ eine $(n \times n)$ -Matrix ist. Dann ist aber AA^{-1} eine $(m \times m)$ -Matrix, und da nach (3.87) $A^{-1}A = AA^{-1}$ gelten soll, folgt $m = n$.

Weiter sei A eine $(n \times n)$ -Matrix und es gelte $AA^{-1} = I$. Dann gilt einerseits $\text{rg}(AA^{-1}) = \text{rg}(I) = n$, andererseits (vergl (3.69)), Seite 82) $\text{rg}(AA^{-1}) \leq \min(\text{rg}(A), \text{rg}(A^{-1}))$, also $\min(\text{rg}(A), \text{rg}(A^{-1})) = n$, woraus $\text{rg}(A) = n$ folgt, d.h. die Zeilen- bzw. Spaltenvektoren von A und A^{-1} sind linear unabhängig. $\square$

Spezialfall: Die Spaltenvektoren der $(n \times n)$ -Matrix A seien paarweise orthonormal. Dann gilt

$$A^{-1} = A', \quad (3.89)$$

d.h. die Inverse von A ist durch die Transponierte von A gegeben.

Beweis: Der Beweis folgt sofort aus dem Begriff der Orthonormalität. $\square$

Beziehung zur Determinante von A : In (3.86) wurde die Beziehung $A(A^*) = |A|I$ aufgestellt. Dann folgt

$$A \frac{(A^*)'}{|A|} = I,$$

so dass

$$\frac{(A^*)'}{|A|} = A^{-1} \quad (3.90)$$

folgt.

Beispiel 3.4 Es sei

$$A = \begin{pmatrix} a & b \\ c & d \end{pmatrix},$$

und gesucht ist

$$A^{-1} = \begin{pmatrix} \alpha & \beta \\ \gamma & \delta \end{pmatrix}.$$

Es ist

$$AA^{-1} = \begin{pmatrix} a\alpha + b\gamma & a\beta + b\delta \\ c\alpha + d\gamma & c\beta + d\delta \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}.$$

Man hat also das Gleichungssystem

$$a\alpha + b\gamma = 1 \Rightarrow \alpha = \frac{1 - b\gamma}{a} \quad (3.91)$$

$$a\beta + b\delta = 0 \Rightarrow \beta = -\frac{b\delta}{a} \quad (3.92)$$

$$c\alpha + d\gamma = 0 \Rightarrow \gamma = -\frac{c\alpha}{d} \quad (3.93)$$

$$c\beta + d\delta = 1 \Rightarrow \delta = \frac{1 - c\beta}{d} \quad (3.94)$$

Durch Einsetzen etwa des Ausdrucks für γ (3.93) in (3.91) etc findet man

$$A^{-1} = \frac{1}{ad - bc} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}. \quad (3.95)$$

Man überprüft durch Nachrechnen, dass in der Tat $AA^{-1} = A^{-1}A = I$ gilt, – vorausgesetzt, dass $ad - bc \neq 0$ ist. Es werde angenommen, dass $ad - bc = 0$ ist. Dann folgt einerseits

$$b = a(d/c), \quad d = c(b/a),$$

andererseits folgt auch

$$d/c = b/a = \lambda,$$

d.h. es gilt

$$\begin{pmatrix} b \\ d \end{pmatrix} = \lambda \begin{pmatrix} a \\ c \end{pmatrix},$$

die Spaltenvektoren von A sind linear abhängig. Man rechnet leicht nach, dass umgekehrt die lineare Abhängigkeit der Spaltenvektoren die Gleichung $ad - bc = 0$ impliziert. Die Voraussetzung $ad - bc \neq 0$ gilt also genau dann, wenn die Spalten- und damit auch die Zeilenvektoren von A linear unabhängig sind. Im Übrigen ist $\det(A) = ad - bc$; eine Lösung existiert also nur, wenn $\det(A) \neq 0$ ist. $\square$

Beispiel 3.5 Es werde der Fall einer Matrix A mit den Zeilen $(1, 3)$ und $(2, 1)$ betrachtet; gesucht ist die zu A inverse Matrix A^{-1} :

$$AA^{-1} = \begin{pmatrix} 1 & 3 \\ 2 & 1 \end{pmatrix} \begin{pmatrix} a & b \\ c & d \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}.$$

Die Spalten von A sind sicher nicht orthogonal: $1 \cdot 3 + 2 \cdot 1 = 5 \neq 0$, aber sie sind linear unabhängig, denn die Spaltenvektoren sind offenbar nicht parallel. Es müssen die Elemente a, b, c und d von A^{-1} bestimmt werden. Man erhält zwei Gleichungssysteme:

$$\begin{array}{lcl} 1 \cdot a + 3 \cdot c = 1 & 1 \cdot b + 3 \cdot d = 0 \\ 2 \cdot a + 1 \cdot c = 0 & 2 \cdot b + 1 \cdot d = 1 \end{array} ,$$

Man findet nun leicht

$$A^{-1} = \begin{pmatrix} -1/5 & 3/5 \\ 2/5 & -1/5 \end{pmatrix} = -\frac{1}{5} \begin{pmatrix} 1 & -3 \\ -2 & 1 \end{pmatrix}.$$

Man rechnet leicht nach, dass $AA^{-1} = A^{-1}A = I$ ist. □

Die lineare Unabhängigkeit der Spalten von A ist in der Tat notwendig für die Existenz einer Inversen. Denn es sei nun

$$AA^{-1} = \begin{pmatrix} 1 & 3 \\ 2 & 6 \end{pmatrix} \begin{pmatrix} a & b \\ c & d \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}.$$

Offenbar sind die Spaltenvektoren von A linear abhängig, da sie sich nur durch einen Faktor unterscheiden und deshalb die gleiche Orientierung haben. Soll A^{-1} existieren, so müssen die Gleichungssysteme

$$\begin{array}{lcl} 1 \cdot a + 3 \cdot c = 1 & & 1 \cdot b + 3 \cdot d = 0 \\ 2 \cdot a + 6 \cdot c = 0 & , & 2 \cdot b + 6 \cdot d = 1 \end{array}$$

eine Lösung haben. Die erste Gleichung im Gleichungssystem links impliziert $a = 1 - 3c$, die zweite Gleichung impliziert $a = -3c$, woraus man $0 = 1$ folgern kann. Beim Gleichungssystem auf der rechten Seite kann man die zweite Gleichung durch 2 teilen, so dass $b + 3d = 1/2$ folgt; zusammen mit der ersten Gleichung folgt nun $0 = 1/2$. Diese Folgerung ist, wie schon die Folgerung $0 = 1$, allenfalls nach Hegel akzeptabel, da nach Ansicht dieses Philosophen das Nichts (die Null) sich negiert und deshalb das Sein (hier $= 1/2$) in Form des Werdens impliziert¹⁸. Nach den Regeln des bloß 'tabellarischen Verstandes' (Hegel) wird man allerdings folgern, dass es keine Lösungen für die beiden Gleichungssysteme gibt und dass daher keine Inverse A^{-1} existiert. Notwendige Voraussetzung für die Existenz der Inversen A^{-1} ist also die lineare Unabhängigkeit der Spalten und damit auch der Zeilen von A ; es kann gezeigt werden, dass für eine quadratische Matrix A diese Bedingung auch hinreichend ist. Der oben eingeführte Ausdruck *singulär* charakterisiert eine Matrix, deren Spalten linear abhängig sind und für die also keine Inverse existiert.

Beispiel 3.6 Es sei R eine (2×2) -Korrelationsmatrix,

$$R = \begin{pmatrix} 1 & r \\ r & 1 \end{pmatrix}. \tag{3.96}$$

Die Inverse R^{-1} ergibt sich aus (3.95)

$$R^{-1} = \begin{pmatrix} \frac{1}{1-r^2} & -\frac{r}{1-r^2} \\ -\frac{r}{1-r^2} & \frac{1}{1-r^2} \end{pmatrix} \tag{3.97}$$

¹⁸Nachzulesen u.a. in Hegels 'Wissenschaft der Logik'.

R^{-1} heißt auch *Präzisionsmatrix*. Der Ausdruck wird klar, wenn man R^{-1} für $r \rightarrow 0$ und $r \rightarrow 1$ bzw $r \rightarrow -1$ betrachtet. Offenbar ist

$$\lim_{r \rightarrow 0} R^{-1} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}, \quad \lim_{r \rightarrow 1} R^{-1} = \begin{pmatrix} \infty & -\infty \\ -\infty & \infty \end{pmatrix} \quad (3.98)$$

Für $r \rightarrow -1$ ändert sich das Vorzeichen der Elemente neben der Diagonalen. Aus der Regressionsrechnung ist bekannt, dass für $y = bx + a + e$

$$r_{xy}^2 = r^2 = 1 - \frac{s_e^2}{s_y^2},$$

d.h. $r \rightarrow 1$, wenn $s_e^2 \rightarrow 0$. Ein kleiner Wert für die Fehlervarianz bedeutet größere Präzision der Vorhersage von y , und dies drückt sich in einem betragsmäßig großen Wert der Elemente von R^{-1} aus. Für $s_e^2 \rightarrow s_y^2$ folgt $r \rightarrow 0$ und y kann nicht aufgrund der x -Werte vorhergesagt werden, d.h. die Präzision geht gegen Null. $\square$

Satz 3.9 Sind A und B $n \times n$ -Matrizen mit existierenden Inversen A^{-1} und B^{-1} , so gilt

$$(AB)^{-1} = B^{-1}A^{-1}. \quad (3.99)$$

Beweis: Nach Voraussetzung existieren A^{-1} und B^{-1} ; dann folgt aus der Assoziativitätsregel (3.34) wegen $AA^{-1} = BB^{-1} = I$ und $(AB)(AB)^{-1} = I$, dass

$$AA^{-1} = AIA^{-1} = A(BB^{-1})A^{-1} = (AB)(B^{-1}A^{-1}) = I,$$

d.h. aber $(AB)^{-1} = B^{-1}A^{-1}$. Damit ist nur gezeigt, dass die *Rechtsinverse* von AB existiert, also die Inverse $(AB)^{-1}$, mit der AB von *rechts* multipliziert wird. Die Argumentation in Bezug auf die Linksinverse ist analog. $\square$

Die Inverse für eine beliebige $(n \times n)$ -Matrix A läßt sich im Prinzip auf die gleiche Weise wie die für eine (2×2) -Matrix finden. Hier soll nur der allgemeine Ausdruck für die Inverse angegeben werden. Es sei A_{ij} eine $(n-1) \times (n-1)$ -Matrix, die aus A entsteht, indem aus A die i -te Zeile und die j -te Spalte gestrichen werden.

Definition 3.9 Die Determinante $|A_{ij}|$ heißt der (i, j) -te Minor von A . Das Element

$$\tilde{a}_{ij} = \frac{(-1)^{i+j}|A_{ij}|}{|A|} \quad (3.100)$$

heißt (i, j) -ter Kofaktor von A . Die $\tilde{a}_{ij}$ können zu einer Matrix $\text{cof}(A)$ zusammengefasst werden; Die Transponierte von $\text{cof}(A)$ heißt Adjunkte von A und wird mit $\text{adj}(A)$ bezeichnet.

Die Inverse von A ist dann durch

$$A^{-1} = \frac{\text{adj}(A)}{|A|} \quad (3.101)$$

gegeben.

Beispiel 3.7 Gegeben sei ein Vektor $\mathbf{w} \in \mathbb{R}^n$ und zwei Basen des V_n , $B = [\mathbf{b}_1, \dots, \mathbf{b}_n]$ und $C = [\mathbf{c}_1, \dots, \mathbf{c}_n]$, d.h. die Basisvektoren $\mathbf{b}_j$ und $\mathbf{c}_j$ werden in den Matrizen $B \in \mathbb{R}^{n,n}$ und $C \in \mathbb{R}^{n,n}$ zusammengefasst. Demnach gilt

$$\mathbf{w} = u_1 \mathbf{b}_1 + u_2 \mathbf{b}_2 + \dots + u_n \mathbf{b}_n = v_1 \mathbf{c}_1 + v_2 \mathbf{c}_2 + \dots + v_n \mathbf{c}_n \quad (3.102)$$

bzw.

$$\mathbf{w} = B\mathbf{u} = C\mathbf{v}. \quad (3.103)$$

Dabei sind $\mathbf{u} = (u_1, \dots, u_n)'$ und $\mathbf{v} = (v_1, \dots, v_n)'$ die jeweiligen Koeffizientenvektoren, die zur Darstellung von $\mathbf{w}$ benötigt werden. Die Frage ist nun, in welcher Beziehung die Vektoren $\mathbf{u}$ und $\mathbf{v}$ zueinander stehen, d.h. wie die Matrix A bestimmt werden kann derart, dass $\mathbf{u} = A\mathbf{v}$ gilt. Da die Spaltenvektoren von B und C linear unabhängig und B und C zudem quadratisch sind kann man folgern, dass die Inversen Matrizen B^{-1} und C^{-1} existieren und man erhält etwa durch Multiplikation von links mit B^{-1}

$$B^{-1}\mathbf{w} = \mathbf{u} = B^{-1}C\mathbf{v} = A\mathbf{v}, \quad (3.104)$$

d.h.

$$A = B^{-1}C \quad (3.105)$$

und damit hat man die Beziehung zwischen den Koeffizientenvektoren $\mathbf{u}$ und $\mathbf{v}$: die Transformationsmatrix A , die $\mathbf{v}$ in $\mathbf{u}$ überführt, ist durch $A = B^{-1}C$ gegeben. Weiter ist

$$A^{-1} = (B^{-1}C)^{-1} = C^{-1}B,$$

so dass auch

$$\mathbf{v} = A^{-1}\mathbf{u}$$

gilt. □

Beispiel 3.8 Es sei A eine $(n \times n)$ -Matrix für die die inverse Matrix A^{-1} existiere. Dann gilt

$$(A^{-1})' = (A')^{-1}. \quad (3.106)$$

Denn $AA^{-1} = I$, und

$$I' = I = (AA^{-1})' = (A^{-1})'A'.$$

Multiplikation von rechts mit $(A')^{-1}$ liefert

$$(A')^{-1} = (A^{-1})'A'(A')^{-1} = (A^{-1})'.$$

□

Im Prinzip läßt sich auch für eine $(m \times n)$ -Matrix A mit etwa $m > n$ eine inverse Matrix $\tilde{A}$ definieren: $\tilde{A}A = I$. Da I nach Definition quadratisch ist, muß $\tilde{A}$ eine $(n \times m)$ -Matrix sein (die Zahl der Spalten von $\tilde{A}$ muß ja stets gleich der Anzahl der Zeilen von A sein, damit die Matrixmultiplikation durchgeführt werden kann, und da A n Spalten hat und I demnach ebenfalls n Spalten und ergo auch n Zeilen haben muß, folgt, dass $\tilde{A}$ n Zeilen haben muß). Dann kann aber nicht auch $A\tilde{A} = \tilde{A}A$ gelten, denn $A\tilde{A}$ ist nun eine $(m \times m)$ -Matrix. $\tilde{A}$ kann also bestenfalls als *Linksinverse* definiert werden, und analog dazu läßt sich unter Umständen eine *Rechtsinverse* definieren. In diesem Abschnitt wird aber nur die in Definition 3.8 spezifizierte Inverse betrachtet.

3.9 Quadratische Formen und Eigenvektoren symmetrischer Matrizen

Es sei A eine beliebige $(m \times n)$ -Matrix, $\mathbf{x}$ sei ein n -dimensionaler Vektor und es gelte $A\mathbf{x} = \mathbf{y}$. $\mathbf{y}$ ist dann m -dimensional. A bildet Vektoren aus einem $\mathbb{R}^n$ auf Vektoren aus einem $\mathbb{R}^m$ ab. Für den Fall $m = n$ (A ist quadratisch) ergeben sich zwei für die Anwendungen in der Multivariaten Statistik wichtige Spezialfälle:

1. A ist eine $(n \times n)$ -Matrix und es gelte $A\mathbf{x} = \mathbf{y}$ für einen beliebigen Vektor $\mathbf{x} \in \mathbb{R}^n$. Dann ist ebenfalls $\mathbf{y} \in \mathbb{R}^n$. Im Allgemeinen unterscheiden sich $\mathbf{x}$ und $\mathbf{y}$ durch ihre Orientierung und durch ihre Länge.

Nun gelte insbesondere $\|\mathbf{x}\| = \|\mathbf{y}\|$, d.h. $\mathbf{x}$ und $\mathbf{y}$ haben identische Längen, d.h. die Transformation A ist *längeninvariant*. $\mathbf{x}$ und $\mathbf{y}$ unterscheiden sich nur durch ihre Orientierungen. A heißt dann auch *Rotationsmatrix*.

2. Für $\mathbf{x} \in \mathbb{R}^n$ gelte $A\mathbf{x} = \mathbf{y} = \lambda\mathbf{x}$. $\mathbf{x}$ und $\mathbf{y}$ haben also dieselbe Orientierung, unterscheiden sich aber im Falle $\lambda \neq 1$ durch ihre Länge. Für eine gegebene Matrix A können die Vektoren $\mathbf{x}$ *nicht* beliebig gewählt werden, die *Orientierungsinvarianz* kann nur für spezielle, für A charakteristische Vektoren $\mathbf{x}$ gelten, die deswegen auch als *charakteristische Vektoren* oder als *Eigenvektoren* von A bezeichnet werden. Dabei kann der wiederum für die Anwendungen sehr wichtige Fall eintreten, dass die zu einer Matrix T zusammengefassten Eigenvektoren einer Matrix A die Eigenschaft einer Rotationsmatrix haben.

3.9.1 Rotationen

Es seien $\mathbf{x}, \mathbf{y}$ zwei n -dimensionale Vektoren und T sei eine Matrix derart, dass $\mathbf{y} = T\mathbf{x}$. T muß eine $(n \times n)$ -Matrix sein, da andernfalls die Vektoren $\mathbf{x}$ und $\mathbf{y}$ nicht beide n -dimensional sein können. T lasse die Länge von $\mathbf{x}$ invariant, so dass sich $\mathbf{y}$ von $\mathbf{x}$ nur in Bezug auf die Orientierung unterscheidet. Dementsprechend soll

$$\|\mathbf{y}\|^2 = \mathbf{y}'\mathbf{y} = \mathbf{x}'T'T\mathbf{x} = \mathbf{x}'\mathbf{x} = \|\mathbf{x}\|^2 \quad (3.107)$$

gelten.

Satz 3.10 *Die Beziehung (3.107) gilt genau dann, wenn die Spaltenvektoren von T orthonormal sind, so dass $T'T = TT' = I$ gilt, wobei I die $(n \times n)$ -Einheitsmatrix ist.*

Beweis: Die Bedingung $T'T = I$ ist sicher hinreichend dafür, dass $\mathbf{x}'T'T\mathbf{x} = \|\mathbf{x}\|^2$ erfüllt ist, denn $\mathbf{x}'T'T\mathbf{x} = \mathbf{x}'I\mathbf{x} = \mathbf{x}'\mathbf{x} = \|\mathbf{x}\|^2$.

Die Beziehung $T'T = I$ ist auch notwendig für die Gültigkeit von (3.107). Um diese Aussage einzusehen, werde $U = T'T$ gesetzt. Nach (3.107) soll $\mathbf{x}'U\mathbf{x} = \mathbf{x}'\mathbf{x}$ für alle $\mathbf{x}$ gelten. Man kann dann die Ableitung von $\mathbf{x}'U\mathbf{x}$ nach $\mathbf{x}$ betrachten:

$$\frac{d(\mathbf{x}'U\mathbf{x})}{d\mathbf{x}} = U\mathbf{x} = I\mathbf{x}, \quad \text{für alle } \mathbf{x},$$

(vergl. (3.193), Seite 120), und nochmalige Ableitung nach $\mathbf{x}$ liefert $U = I_n$ (vergl. (3.190), Seite 120), d.h. $T'T = I$. Dies bedeutet, dass die Spaltenvektoren von T orthonormal sind. $\square$

Satz 3.11 *Eine Rotation läßt die Skalarprodukte zwischen den rotierten Vektoren invariant.*

Beweis: Für $\mathbf{u} = T\mathbf{x}$, $\mathbf{v} = T\mathbf{y}$ folgt sofort

$$\mathbf{u}'\mathbf{v} = \mathbf{x}'T'T\mathbf{y} = \mathbf{x}'\mathbf{y}, \quad T'T = I. \quad (3.108)$$

$\square$

Folgerung: Aus $T'T = I$, I die Einheitsmatrix, folgt

$$T' = T^{-1}, \quad (3.109)$$

d.h. die Transponierte T' ist gleich der inversen Matrix von T .

Satz 3.12 *Die Spaltenvektoren von T seien orthonormal. Dann sind auch die Zeilenvektoren von T orthonormal,*

$$T'T = I \Rightarrow TT' = I. \quad (3.110)$$

Beweis: Es gilt $T = TI = T(T'T) = (TT')T = AT$, mit $A = TT'$. Dann

$$T - AT = 0 \Rightarrow (T - AT)' = T' - T'A' = T'(I - A) = 0,$$

wegen $A' = A$, und 0 die Nullmatrix. Da der Zeilenrang stets gleich dem Spaltenrang ist und T quadratisch ist, folgt, dass die Zeilenvektoren von T – also die Spaltenvektoren von T' – linear unabhängig sind. Es sei $B = I - A$ und es sei

$\mathbf{b}_j$ der j -te Spaltenvektor von B ; dann gilt also $T'\mathbf{b}_j = \vec{0}_j$, $\vec{0}_j$ der j -te Spaltenvektor der Nullmatrix 0 , der hier als Linearkombination der Spaltenvektoren von T' dargestellt wird. Wegen der linearen Unabhängigkeit der Spaltenvektoren von T' folgt dann $\mathbf{b}_j = \vec{0}$ für alle j , d.h. $I - A = 0$, also $A = TT' = I$, so dass T' orthonormal ist. $\square$

Anmerkung: Setzt man den Begriff der Determinante einer Matrix voraus, so läßt sich auf der Basis von Satz 3.5, Seite 86, ein alternativer Beweis finden: A sei orthonormal. Nach Satz 3.5, A5 gilt $|I| = 1$, denn I ist sicherlich eine obere Dreiecksmatrix. Nach A7 gilt $|A| = |A'|$, und nach A6 gilt

$$|A'A| = |A'| |A| = |AA'|,$$

mithin folgt $|A'A| = |I| = |AA'| = 1$, d.h. $AA' = I$. $\square$

T läßt sich durch trigonometrische Betrachtungen zur Rotation (etwa von Koordinatensystemen) herleiten; im 2-dimensionalen Fall erhält man für T den Ausdruck

$$T = \begin{pmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{pmatrix}, \quad (3.111)$$

wobei θ der Rotationswinkel ist. Eine gegebene $(n \times n)$ -Rotationsmatrix T rotiert alle n -dimensionalen Vektoren $\mathbf{y}$ um einen bestimmten, fixen Winkel θ , so dass man auch $T(\theta)$ schreiben könnte, um diesen Sachverhalt auszudrücken. Im Folgenden wird von dieser Beziehung zwischen θ und T kein Gebrauch gemacht, weil θ nicht explizit in die Betrachtungen eingeht.

3.9.2 Quadratische Formen und Eigenvektoren

Übersicht: Es wird zunächst gezeigt, dass bestimmten, d.h. positiv semidefiniten symmetrischen Matrizen M Ellipsoide zugeordnet werden können; jedem Fall (z.B. einer Zeile in einer Matrix $(m \times n)$ -Matrix X mit $M = X'X$ mit m Fällen und n Variablen) entspricht ein Punkt in einem n -dimensionalen Raum, und jedem dieser Punkte entspricht ein Ellipsoid, auf dem der Punkt liegt. Alle Ellipsoide haben dieselbe Orientierung. Die Orientierung ist durch die Eigenvektoren von M gegeben. $\square$

Definition 3.10 Es sei M eine symmetrische $(n \times n)$ -Matrix und $\mathbf{x} \in \mathbb{R}^n$, und es gelte

$$Q_M(\mathbf{x}) = \mathbf{x}'M\mathbf{x} \quad (3.112)$$

Dann heißt $Q_M(\mathbf{x})$ quadratische Form.

Definition 3.11 Es sei $M \in \mathbb{R}^{n,n}$ eine symmetrische Matrix. M heißt positiv semidefinit, wenn $\mathbf{x}'M\mathbf{x} \geq 0$ für alle Vektoren $\mathbf{x} \in \mathbb{R}^n$ gilt. M heißt negativ semidefinit, wenn für alle $\mathbf{x} \in \mathbb{R}^n$ die Beziehung $\mathbf{x}'M\mathbf{x} \leq 0$ gilt. M heißt positiv definit bzw. elliptisch, wenn für alle $\mathbf{x} \in \mathbb{R}^n$ $\mathbf{x}'M\mathbf{x} > 0$, und negativ definit bzw. hyperbolisch, wenn $\mathbf{x}'M\mathbf{x} < 0$ gilt.

Satz 3.13 Es sei $X \in \mathbb{R}^{m,n}$. Dann ist $M = X'X$ positiv semidefinit.

Beweis: $M = X'X$ ist eine symmetrische $(n \times n)$ -Matrix. $\mathbf{x}$ sei ein n -dimensionaler Vektor, und es sei $M\mathbf{x} = \mathbf{y} \in \mathbb{R}^n$. Dann ist $\mathbf{x}'M\mathbf{x} = \mathbf{x}'X'X\mathbf{x} = \mathbf{y}'\mathbf{y} \geq 0$, d.h. M ist positiv semidefinit. $\square$

Anmerkung: Insbesondere ist jede Varianz-Kovarianz-Matrix S positiv semidefinit. Nach Gleichung (3.49), Seite 75 kann S in der Form

$$S = \frac{1}{m}X'X - \bar{\mathbf{x}}\bar{\mathbf{x}}' = \frac{1}{m}(X'X - \frac{1}{m}X'\bar{\mathbf{1}}\bar{\mathbf{1}}'X).$$

geschrieben werden, wobei H die Zentrierungsmatrix ist. Auf Seite 75 wurde gezeigt, dass H symmetrisch und idempotent ist, d.h. es gilt $HH = H'H = H^2 = H$, d.h. $mS = X'HX = X'H'HX$, und für einen beliebigen, n -dimensionalen Vektor $\mathbf{x} \neq \vec{0}$ folgt mit $\mathbf{y} = HX\mathbf{x}$

$$\mathbf{x}'S\mathbf{x} = \mathbf{x}'X'H'HX\mathbf{x} = \|\mathbf{y}\|^2 \geq 0,$$

d.h. S ist positiv semidefinit. $\square$

Satz 3.14 Es sei $M \in \mathbb{R}^{n,n}$ eine symmetrische, positiv semidefinite Matrix. Dann definiert die Menge $\mathcal{E}_x = \{\mathbf{x}'M\mathbf{x} = k, \mathbf{x} \in \mathbb{R}^n\}$ ein n -dimensionales Ellipsoid, wobei die Anfangspunkte der $\mathbf{x}$ im Nullpunkt des Koordinatensystems liegen und die Endpunkte auf dem jeweiligen Ellipsoid.

Beweis: Die Aussage folgt sofort aus der Definition von $Q_M(\mathbf{x})$: multipliziert man (3.112) aus, so erhält man

$$\mathbf{x}'M\mathbf{x} = \sum_{i=1}^n m_{ii}x_i^2 + 2 \sum_{i < j} m_{ij}x_i x_j = k \quad (3.113)$$

Für $\mathbf{x}'M\mathbf{x} = k > 0$ definiert $\mathbf{x}'M\mathbf{x}$ ein n -dimensionales Ellipsoid. $\square$

Der Ausdruck 'quadratische Form' ergibt sich aus dem Sachverhalt, dass die Summe der Exponenten der Komponenten x_i stets gleich 2 ist. Für den Spezialfall $n = 2$ hat man

$$\mathbf{x}'M\mathbf{x} = m_{11}x_1^2 + m_{22}x_2^2 + 2m_{12}x_1x_2 = k. \quad (3.114)$$

Die Menge der 2-dimensionalen Vektoren $\mathbf{x} = (x_1, x_2)'$, die dieser Gleichung genügen, definiert eine Ellipse (s. a. Satz 3.19, Seite 102).

Spezialfall: Insbesondere sei $M = \Lambda = \text{diag}(\lambda_1, \dots, \lambda_n)$ eine Diagonalmatrix. Dann sind die Ellipsoide $\mathbf{x}'\Lambda\mathbf{x} = k$ *achsenparallel*, d.h. die Hauptachsen der Ellipsoide sind parallel zu den Achsen des Koordinatensystems; diese Aussage folgt sofort aus (3.113) bzw. (3.114), denn für $M = \Lambda$ sind alle $m_{ij} = 0$ für $i \neq j$. In Beispiel 3.11, Seite 106 wird diese Aussage noch einmal elaboriert.

Es seien nun $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$ Vektoren und M sei eine symmetrische $(n \times n)$ -Matrix, und es gelte

$$\mathbf{x}'M\mathbf{x} = \mathbf{y}'\Lambda\mathbf{y} = k > 0, \quad \Lambda = \text{diag}(\lambda_1, \dots, \lambda_n). \quad (3.115)$$

$\mathcal{E}_x = \{\mathbf{x} | \mathbf{x}'M\mathbf{x} = k > 0\}$ und $\mathcal{E}_y = \{\mathbf{y} | \mathbf{y}'\Lambda\mathbf{y} = k > 0\}$ sind Ellipsoide, $\mathcal{E}_y$ ist insbesondere achsenparallel. Weiter gelte $\mathbf{x} = T\mathbf{y}$, wobei T eine Rotation repräsentiere. Offenbar gilt für alle $\mathbf{x} \in \mathcal{E}_x$ und alle $\mathbf{y} \in \mathcal{E}_y$

$$\mathbf{x}'M\mathbf{x} = \mathbf{y}'TMT\mathbf{y} = \mathbf{y}'\Lambda\mathbf{y} = k,$$

woraus

$$T'MT = \Lambda \quad (3.116)$$

folgt (eine ausführliche Begründung dieser Folgerung findet man in Beispiel 3.13, Seite 121). Denn $\mathcal{E}_y$ ist ein achsenparalleles Ellipsoid, so dass auch $\mathbf{y}'TMT\mathbf{y} = k$ ein achsenparalleles Ellipsoid sein muß, d.h. $T'MT = D$ muß eine Diagonalmatrix sein, $D = \text{diag}(d_1, \dots, d_n)$. Dann folgt aber

$$\mathbf{y}'\Lambda\mathbf{y} = \mathbf{y}'D\mathbf{y} = k,$$

d.h.

$$\sum_{k=1}^n \lambda_k y_k^2 = \sum_{k=1}^n d_k y_k^2.$$

Differenziert man beide Seiten nach y_k , so erhält man $\lambda_k = d_k$, d.h. $\Lambda = D$, und das ist (3.116). $\square$

Da T als Rotationsmatrix angenommen wurde, folgt, dass T orthonormal ist. Deshalb folgt durch Multiplikation der Gleichung (3.116) von links mit T die Gleichung

$$MT = T\Lambda. \quad (3.117)$$

Diese Gleichung besagt, dass die Spaltenvektoren von T durch M so transformiert werden, dass sich nur ihre Länge, nicht aber ihre Orientierung verändert. Diese Aussage gilt natürlich nicht für beliebige Vektoren $\mathbf{x}$, sondern nur für spezielle Vektoren $\mathbf{t}$, die charakteristisch für die Matrix M sind.

Definition 3.12 *Es sei M eine beliebige $(n \times n)$ -Matrix und $\mathbf{t} \in \mathbb{R}^n$ sei ein Vektor, der der Beziehung*

$$M\mathbf{t} = \lambda\mathbf{t}, \quad \mathbf{t} \neq \vec{0}, \quad (3.118)$$

genügt. Dann heißt $\mathbf{t}$ Eigenvektor von M und λ heißt der zu $\mathbf{t}$ gehörende Eigenwert von M .

Die Gleichung (3.117) besagt also, dass alle Spaltenvektoren $\mathbf{t}_k$, $k = 1, \dots, n$, von T Eigenvektoren von M sind, und die Diagonalmatrix Λ enthält in der Diagonalen die zugehörigen Eigenwerte von M .

Bemerkungen: In der Definition wurde *nicht* vorausgesetzt, dass M symmetrisch ist, d.h. Eigenvektoren können auch für nicht-symmetrische Matrizen existieren. Die folgenden Betrachtungen beschränken sich aber auf symmetrische Matrizen M . Insbesondere können Eigenwerte komplexe Zahlen sein: $\lambda = x + iy$, $i = \sqrt{-1}$, und die Eigenvektoren können komplexe Komponenten haben. $\square$

Eine Rotationsmatrix S ist orthonormal, und wenn in einem bestimmten Kontext gefolgert wird, dass $S = T$ auch eine Matrix von Eigenvektoren ist, so folgt, dass T ebenfalls orthonormal ist. Nun sei umgekehrt bekannt, dass T eine Matrix von Eigenvektoren von $M' = M$ ist. Die Frage ist, ob nun auch folgt, dass T orthonormal ist, – es ist ja denkbar, dass man nicht über eine Rotation auf die Eigenschaft der Spalten von T , Eigenvektoren zu sein, gekommen ist, und vielleicht gibt es auch nicht-orthogonale Eigenvektoren von M . Dazu wird der folgende Satz bewiesen:

Satz 3.15 $M \in \mathbb{R}^{n,n}$ sei symmetrisch und habe die Eigenvektoren $\mathbf{t}_1, \dots, \mathbf{t}_n$ mit den korrespondierenden Eigenwerten $\lambda_1, \dots, \lambda_n$. Die Eigenwerte sind stets reell, d.h. $\lambda_k \in \mathbb{R}$ für alle k . Sind $\mathbf{t}_j$ und $\mathbf{t}_k$ mit zugehörigen Eigenwerten $\lambda_j \neq \lambda_k$ irgendzwei Eigenvektoren von M , so sind $\mathbf{t}_j$ und $\mathbf{t}_k$ orthogonal, d.h. es gilt

$$\mathbf{t}'_j \mathbf{t}_k = \begin{cases} 0, & j = k \\ \|\mathbf{t}_j\| \neq 0, & j \neq k \end{cases} \quad (3.119)$$

Beweis: Es sei $M' = M$ und es gelte $M\mathbf{t} = \lambda\mathbf{t}$. Es werde angenommen, dass $\lambda \in \mathbb{C}$ und $\mathbf{v} \in \mathbb{C}^n$, $\mathbb{C}$ die Menge der komplexen Zahlen. Gilt $z = x + iy \in \mathbb{C}$ so heißt $\bar{z} = x - iy$ die zu z konjugiert komplexe Zahl und es gilt $z\bar{z} = (x + ia)(x - iy) = x^2 + y^2$, da ja $i^2 = -1$. Dann folgt

$$\bar{\lambda} \bar{\mathbf{v}}' \mathbf{v} = \bar{\lambda} \|\mathbf{v}\|^2 = (M \bar{\mathbf{v}}') \mathbf{v} = \bar{\mathbf{v}}' M \mathbf{v} = \bar{\mathbf{v}}' \lambda \mathbf{v} = \lambda \|\mathbf{v}\|^2,$$

d.h. es gilt $\bar{\lambda} = \lambda \in \mathbb{R}$.

Für irgendzwei Eigenvektoren $\mathbf{t}_j$ und $\mathbf{t}_k$ gelte $\lambda_j \neq \lambda_k$. Dann sind $\mathbf{t}_j$ und $\mathbf{t}_k$ orthogonal. Denn dann gilt

$$M\mathbf{t}_j = \lambda_j \mathbf{t}_j \quad (3.120)$$

$$M\mathbf{t}_k = \lambda_k \mathbf{t}_k \quad (3.121)$$

Die Gleichung (3.120) werde von links mit $\mathbf{t}'_k$, die Gleichung (3.121) von links mit $\mathbf{t}_j$ multipliziert. Es entstehen die Gleichungen

$$\mathbf{t}'_k M \mathbf{t}_j = \lambda_j \mathbf{t}'_k \mathbf{t}_j \quad (3.122)$$

$$\mathbf{t}'_j M \mathbf{t}_k = \lambda_k \mathbf{t}'_j \mathbf{t}_k. \quad (3.123)$$

Nun ist einerseits $\mathbf{t}'_j \mathbf{t}_k = \mathbf{t}'_k \mathbf{t}_j$, und andererseits $(\mathbf{t}'_k M \mathbf{t}_j)' = \mathbf{t}'_j M \mathbf{t}_k$, da ja $M' = M$. Subtrahiert man also die zweite Gleichung von der ersten, ergibt sich

$$0 = (\lambda_j - \lambda_k) \mathbf{t}'_j \mathbf{t}_k,$$

woraus wegen $\lambda_j - \lambda_k \neq 0$ die Behauptung $\mathbf{t}'_j \mathbf{t}_k = 0$ folgt. $\square$

Anmerkung: In (3.119) ist nicht gefordert worden, dass $\|\mathbf{t}_j\| = 1$ ist; $M\mathbf{t}_j = \lambda_j \mathbf{t}_j$ bedeutet, dass sich die Längen der Vektoren $M\mathbf{t}_j$ und $\mathbf{t}_j$ um den Faktor λ_j unterscheiden, unabhängig von der Länge von $\mathbf{t}_j$. Insofern ist die Länge eines Eigenvektors irrelevant und deswegen kann $\|\mathbf{t}_j\| = 1$ gesetzt werden. Ist bereits bekannt, dass T auch eine Rotationsmatrix ist, so wird die Normiertheit der $\mathbf{t}_j$ gewissermaßen gleich mitgeliefert. $\square$

Die Frage ist nun, welche Aussage über die Eigenvektoren einer symmetrischen Matrix gemacht werden kann, wenn nicht alle Eigenwerte voneinander verschieden sind. Der folgende Satz macht hierüber eine Aussage.

Satz 3.16 *Es sei λ_j ein Eigenwert der symmetrischen Matrix M mit der Mehrfachheit m , d.h. es gelte $\lambda_j = \lambda_{j+1} = \dots = \lambda_{j+m}$. Dann existieren m orthogonale, zu λ_j korrespondierende Eigenvektoren.*

Beweis: Vergl. Abschnitt 5, Seite 178. $\square$

Die Orthonormalität von T bedeutet, dass man aus $MT = T\Lambda$ (Gleichung (3.117)) durch Multiplikation von rechts mit T' die Beziehung

$$M = T\Lambda T' = \sum_{k=1}^n \lambda_k \mathbf{t}_k \mathbf{t}'_k. \quad (3.124)$$

erhält. Der Ausdruck $\sum_k \lambda_k \mathbf{t}_k \mathbf{t}'_k$ drückt $T\Lambda T'$ über die dyadischen Produkte $\mathbf{t}_k \mathbf{t}'_k$ aus und erweist bei bestimmten Betrachtungen als nützlich. Man macht sich leicht klar, wie dieser Ausdruck zustande kommt. Es ist ja $T = [\mathbf{t}_1, \dots, \mathbf{t}_n]$ und $T\Lambda = [\lambda_1 \mathbf{t}_1, \dots, \lambda_n \mathbf{t}_n]$, so dass

$$T\Lambda T' = [\lambda_1 \mathbf{t}_1, \dots, \lambda_n \mathbf{t}_n] \begin{bmatrix} \mathbf{t}'_1 \\ \mathbf{t}'_2 \\ \vdots \\ \mathbf{t}'_n \end{bmatrix} = \lambda_1 \mathbf{t}_1 \mathbf{t}'_1 + \lambda_2 \mathbf{t}_2 \mathbf{t}'_2 + \dots + \lambda_n \mathbf{t}_n \mathbf{t}'_n = \sum_{k=1}^n \lambda_k \mathbf{t}_k \mathbf{t}'_k.$$

Definition 3.13 *Die Darstellung (3.124) von M heißt Spektraldarstellung von M .*

Satz 3.17 *Es sei M eine reelle, symmetrische $(n \times n)$ -Matrix. M ist positiv semidefinit dann und nur dann, wenn die Eigenwerte λ_j größer als bzw. mindestens gleich Null sind.*

Beweis: $M = T'\Lambda T$, T die orthonormalen Eigenvektoren von M , impliziert $T'MT = \Lambda$. Weiter sei $Q = \mathbf{x}'M\mathbf{x} \in \mathbb{R}$, $\mathbf{x}$ ein beliebiger n -dimensionaler Vektor. Für einen geeignet gewählten Vektor $\mathbf{y}$ ist $\mathbf{x} = T\mathbf{y}$, so dass

$$Q = \mathbf{y}'T'MT\mathbf{y} = \mathbf{y}'\Lambda\mathbf{y} = \sum_{i=1}^n \lambda_i y_i^2$$

Da $\mathbf{x}$ beliebig gewählt werden kann, kann insbesondere $\mathbf{x} = T\mathbf{e}_k$ gewählt werden, also $\mathbf{y}_k = \mathbf{e}_k$ der k -te Einheitsvektor, $k = 1, \dots, n$. Dann ist $Q = \lambda_k$, und $Q \geq 0$ genau dann, wenn $\lambda_j \geq 0$.

Man zeigt auf analoge Weise, dass für eine negativ (semi-)definite Matrix $\lambda_j \leq 0$ für alle j gilt. $\square$

Satz 3.18 *Es sei M eine reelle, symmetrische $(n \times n)$ -Matrix. Dann ist der Rang von M gleich der Anzahl von Null verschiedenen Eigenwerte.*

Beweis: Der Beweis macht implizit von Satz 3.3, Seite 82, Gebrauch. Es gilt $M = T\Lambda T'$, und es seien $r \leq n$ Eigenwerte ungleich Null. Dann enthält Λ $n - r$ Spalten (und Zeilen), die nur Nullen enthalten. Die Spaltenvektoren von M sind Linearkombinationen der Spalten von $T\Lambda$, so dass für den j -ten Spaltenvektor $\mathbf{v}_j$ von M

$$\mathbf{v}_j = \lambda_1 t_{j1} \mathbf{t}_1 + \dots + \lambda_r t_{jr} \mathbf{t}_r + \underbrace{0 + \dots + 0}_{n-r}$$

gilt. Es genügt demnach,

$$M = T_r \Lambda_r T_r' \quad (3.125)$$

zu schreiben, wobei T_r die Matrix der Eigenvektoren ist, die zu von Null verschiedenen Eigenwerten korrespondieren, die in der Matrix Λ_r zusammengefaßt werden. Die Matrix $T_r \Lambda_r$ besteht aus den Spaltenvektoren $\lambda_j \mathbf{t}_j$, $j = 1, \dots, r$, die orthogonal und damit linear unabhängig sind, mithin hat $T_r \Lambda_r$ den Rang r . (3.125) bedeutet, dass die Spaltenvektoren von M sich als Linearkombinationen der $\lambda_j \mathbf{t}_j$ darstellen lassen, d.h. die Spaltenvektoren sind Elemente der linearen Hülle $\mathcal{L}(\lambda_1 \mathbf{t}_1, \dots, \lambda_r \mathbf{t}_r)$, und somit ist $\text{rg}(M) \leq r$. Aber T_r ist orthonormal, so dass aus (3.125) $MT_r = T_r \Lambda_r$ folgt. Dies heißt aber, dass sich die Spaltenvektoren $\lambda_j \mathbf{t}_j$ von $T_r \Lambda_r$ als Linearkombinationen der Spalten von M darstellen lassen, d.h. sie liegen in der linearen Hülle $\mathcal{L}(M)$ von M . Dies bedeutet, dass $r = \text{rg}(T_r \Lambda_r) \leq \text{rg}(M)$ sein muß. Es muß also $\text{rg}(M) \leq r$ und $\text{rg}(M) \geq r$ gelten, so dass $\text{rg}(M) = r$ folgt. $\square$

Der Satz 3.17 sagt noch wenig aus über die Eigenschaften einer symmetrischen Matrix (es sind stets reelle Matrizen gemeint), die nur positive Eigenwerte implizieren (oder nur negative). Der folgende Satz gibt weitere Auskunft.

Satz 3.19 *Es sei M eine symmetrische $(n \times n)$ -Matrix vom Rang $r \leq n$. Dann ist M genau dann positiv semidefinit, wenn eine $(n \times r)$ -Matrix G existiert derart, dass*

$$M = GG'. \quad (3.126)$$

Beweis: (1) $\Rightarrow$: Es gelte $M = GG'$. Dann folgt

$$\mathbf{x}'GG'\mathbf{x} = (G\mathbf{x})'G\mathbf{x} = \|G\mathbf{x}\|^2 \geq 0,$$

so dass M positiv semidefinit ist.

(2) $\Leftarrow$: Aus der Symmetrie von M folgt die Existenz der Matrizen T (orthonormal) und $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_j)$, $\lambda_j \geq 0$ für alle j (s. Satz 3.17), mit $M = T\Lambda T'$. Es sei

$$\Lambda^{1/2} = \text{diag}(\sqrt{\lambda_1}, \dots, \sqrt{\lambda_r}, \underbrace{0, \dots, 0}_{n-r}).$$

Dann kann man

$$M = T\Lambda^{1/2}\Lambda^{1/2}T' = (T\Lambda^{1/2})(T\Lambda^{1/2})'$$

schreiben. Streicht man in $T\Lambda^{1/2}$ alle Spalten, die nur Nullen enthalten, so erhält man eine Matrix $G = T_r\Lambda_r^{1/2}$ und M ist in der Form $M = GG'$ darstellbar. $\square$

Der Satz 3.19 spezifiziert die Bedingungen, die eine symmetrische Matrix M erfüllen muß, um eine Ellipse bzw. ein Ellipsoid zu definieren.

Beispiel 3.9 Die Korrelation zwischen zwei zufälligen Veränderlichen sei r und die zugehörige Korrelationsmatrix ist durch

$$R = \begin{pmatrix} 1 & r \\ r & 1 \end{pmatrix} \quad (3.127)$$

gegeben. Gesucht sind die Eigenvektoren und die zugehörigen Eigenwerte von R .

Es sei λ ein Eigenwert und $\mathbf{v} = (v_1, v_2)'$ der zugehörige Eigenvektor. Dann gilt

$$\begin{pmatrix} 1 & r \\ r & 1 \end{pmatrix} \begin{pmatrix} v_1 \\ v_2 \end{pmatrix} = \lambda \begin{pmatrix} v_1 \\ v_2 \end{pmatrix}.$$

Ausgeschrieben entsteht das System von Gleichungen

$$v_1 + rv_2 = \lambda v_1 \quad (3.128)$$

$$rv_1 + v_2 = \lambda v_2, \quad (3.129)$$

woraus sich

$$v_1(1 - \lambda) = -rv_2 \quad (3.130)$$

$$v_2(1 - \lambda) = -rv_1 \quad (3.131)$$

ergibt, und daraus folgt

$$\frac{v_1}{v_2} = \frac{v_2}{v_1},$$

also $v_1^2 = v_2^2$. Daraus ergeben sich die Lösungen

$$v_1 = v_2, \quad -v_1 = v_2, \quad (3.132)$$

so dass man die Eigenvektoren

$$\mathbf{v}_1 = \begin{pmatrix} v_1 \\ v_1 \end{pmatrix}, \quad \mathbf{v}_2 = \begin{pmatrix} v_1 \\ -v_1 \end{pmatrix} \quad (3.133)$$

erhält. Die Lösung $v_1 = v_2$ in (3.128) eingesetzt ergibt

$$v_1(1 - \lambda) = -rv_1,$$

und man erhält $\lambda = \lambda_1 = 1 + r$. Setzt man die Lösung $v_2 = -v_1$ in (3.128) oder (3.131) ein, erhält man

$$v_1(1 - \lambda) = rv_1, \quad v_2(1 - \lambda) = rv_2,$$

woraus in jedem Fall $\lambda = \lambda_2 = 1 - r$ folgt, d.h.

$$\lambda = \begin{cases} \lambda_1 = 1 + r \\ \lambda_2 = 1 - r \end{cases} \quad (3.134)$$

Man überprüfe, ob die Gleichungen $R\mathbf{v}_1 = \lambda_1\mathbf{v}_1$ und $R\mathbf{v}_2 = \lambda_2\mathbf{v}_2$ tatsächlich gelten, und ob $\mathbf{v}_1$ und $\mathbf{v}_2$ orthogonal sind. $\mathbf{v}_1 = (v, v)'$ bedeutet, dass $\mathbf{v}_1$ auf der Hauptdiagonalen des Koordinatensystems liegt, das durch die standardisierten Variablen definiert wird, und $\mathbf{v}_2$ liegt auf der Nebendiagonalen. Für $r = 0$ ist $R = I$ die Einheitsmatrix und es folgt $\lambda_1 = \lambda_2 = 1$. $\square$

In den folgenden Abschnitten wird die Rolle von Rotationsmatrizen und Matrizen von Eigenvektoren elaboriert.

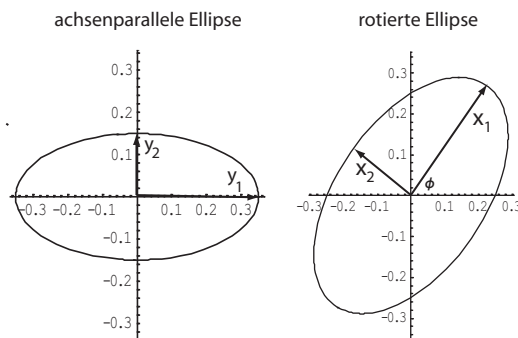
Zur Bedeutung der Eigenvektoren positiv semidefiniter Matrizen: Es sei insbesondere $\mathbf{y}_1 = y\mathbf{e}_1$, $\mathbf{e}_1 = (1, 0, \dots, 0)'$, so dass $\|\mathbf{y}_1\| = y_1\|\mathbf{e}_1\| = y_1$. $\mathbf{y}_1$ definiert dann die erste Halbachse des durch Λ definierten achsenparallelen Ellipsoids. Dann folgt

$$\mathbf{x}_1 = T\mathbf{y}_1 = y_1T\mathbf{e}_1 = y \begin{pmatrix} t_{11} \\ t_{21} \\ \vdots \\ t_{n1} \end{pmatrix} = y\mathbf{t}_1, \quad y_1 \in \mathbb{R}, \quad (3.135)$$

d.h. $\mathbf{x}_1$ ist proportional zum ersten Eigenvektor von M , der in der ersten Spalte von T steht. $\mathbf{t}_1$ definiert die Orientierung der ersten Hauptachse des durch M definierten Ellipsoids. Da die $\mathbf{t}_k$ orthogonal sind, definieren die restlichen Vektoren $\mathbf{t}_k$, $k \neq 1$, die Orientierungen der restlichen Hauptachsen des Ellipsoids.

T rotiert alle Vektoren $\mathbf{y}$, für die $\mathbf{y}'\Lambda\mathbf{y} = k$ gilt. $\mathbf{x}_1$ definiert dann die erste Halbachse des Ellipsoids $\mathbf{x}'M\mathbf{x} = k$. Nach (3.135) ist die Orientierung dieser Halbachse durch den ersten Eigenvektor $\mathbf{t}_1$ von M gegeben. Analoge Interpretationen

Abbildung 7: Ellipsen in verschiedenen Orientierungen und ihre skalierten Eigenvektoren



ergeben sich für die übrigen Halbachsen des durch M definierten Ellipsoids:

$$\mathbf{x}_j = T y_j \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix} = y_j \begin{pmatrix} t_{1j} \\ t_{2j} \\ \vdots \\ t_{nj} \end{pmatrix}, \quad \mathbf{x}_j \in \mathcal{E}_x, \quad \mathbf{y}_j = y_j \mathbf{e}_j \in \mathcal{E}_y \quad (3.136)$$

y_j ist die Länge der jeweiligen Halbachse.

Beispiel 3.10 Ellipsen und Punktekongfigurationen: Es werden Messwertpaare für zwei Variablen erhoben (z.B. werden bei einer Anzahl m von Personen ("Fälle") die Körpergröße (Variable 1) und das Körpergewicht (Variable 2) gemessen. Jede Person wird durch einen Punkt im durch die Variablen 1 und 2 definierten Koordinatensystem repräsentiert, s. Abbildung 10, links. Man kann die Regressionsgeraden für die Vorhersage von Variable 2 durch Variable 1 und umgekehrt bestimmen. Die Steigungen

$$b_{12} = \frac{Kov(V_1, V_2)}{s_1^2}, \quad b_{21} = \frac{Kov(V_1, V_2)}{s_2^2}$$

sind verschieden, wenn die Varianzen s_1^2 und s_2^2 verschieden sind. Die Matrix der Korrelationen ist durch

$$R = \begin{pmatrix} 1 & r_{12} \\ r_{21} & 1 \end{pmatrix}$$

gegeben, wobei natürlich $r_{12} = r_{21}$ ist; in Beispiel 3.9, Seite 103, werden die Eigenvektoren und Eigenwerte dieser Matrix hergeleitet. R definiert eine Menge

von Ellipsen mit identischer Orientierung; für jeden Punkt, also für jeden Fall, kann eine Ellipse bestimmt werden. Ist $\tilde{\mathbf{x}}_i = (x_{i1}, x_{i2})'$ der i -te Zeilenvektor von X , d.h. der i -te Spaltenvektor von X' , so gilt

$$\tilde{\mathbf{x}}_i' R \tilde{\mathbf{x}}_i = k_i,$$

und die Menge der 2-dimensionalen Vektoren $\mathbf{x}$, die der Bedingung $\mathbf{x}' R \mathbf{x} = k_i$ genügen, ist die zum Fall i korrespondierende Ellipse. Für eine Teilmenge der Fälle sind die korrespondierenden Ellipsen in Abbildung 10 eingezeichnet worden.

Die Hauptachsen dieser Ellipse können als neues Koordinatensystem gewählt werden; die Koordinaten der Fälle auf diesen Achsen sind durch die Projektionen der Fälle auf die Hauptachsen gegeben. Ist $\tilde{\mathbf{y}}_i$ der Vektor, der den i -ten Fall im durch die Hauptachsen gegebenen Koordinatensystem repräsentiert, so ist die Beziehung dieses Vektors zum korrespondierenden Vektor im ursprünglichen Koordinatensystem durch $\tilde{\mathbf{x}}_i = T \tilde{\mathbf{y}}_i$ gegeben, und T ist die Matrix der Eigenvektoren von $X'X$. Abbildung 10, rechts, zeigt die Konfiguration im Hauptachsensystem. Offenbar ist die Kovarianz zwischen den Hauptachsen gleich Null. Fasst man die Hauptachsen als Repräsentation latenter Variablen auf, so kann man sagen, dass diese latenten Variablen unkorreliert sind. Die Fälle unterscheiden sich hinsichtlich der ersten Hauptachse maximal. Man könnte vermuten, dass sie eine sllgemeine Größendimension abbildet. Die zweite Hauptachse könnte eine von der Größe unabhängige Variable, etwa die Intensität des Stoffwechsels repräsentieren. Diese Interpretation setzt voraus, dass die Variation der Fälle in Bezug auf diese Achse nicht nur zufällige Effekte reflektiert. $\square$

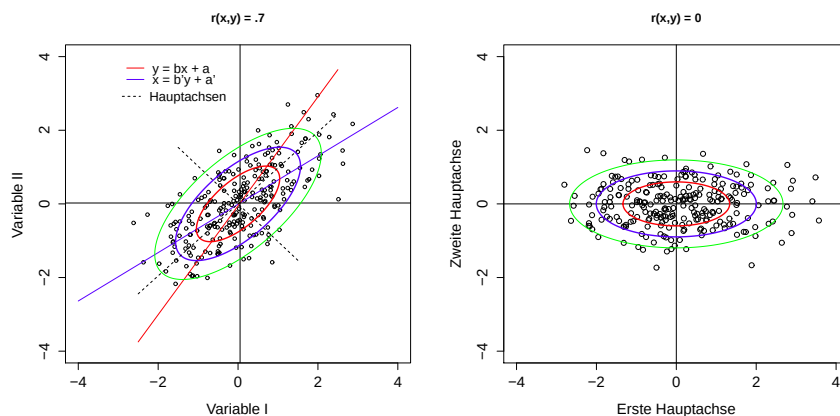
Beispiel 3.11 Eigenvektoren einer Diagonalmatrix Auf Seite 98 wurden quadratische Formen $\mathbf{x}' M \mathbf{x} = k$ für den Fall, dass M eine Diagonalmatrix ist, betrachtet; es wurde gesagt, dass die korrespondierenden Ellipsoide dann achsenparallel seien.

Es sei $M = \text{diag}(\lambda_1, \dots, \lambda_n)$ mit $\lambda_k > 0$ für $k = 1, \dots, n$, so dass M positiv definit ist. Die Eigenvektoren von M sind dann durch die Einheitsvektoren $\mathbf{e}_j$, $j = 1, \dots, n$ mit den zugehörigen Eigenwerten $\lambda = \lambda_k$ gegeben. Denn

$$\begin{pmatrix} \lambda_1 & 0 & 0 & \dots & 0 \\ 0 & \lambda_2 & 0 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \dots & \lambda_n \end{pmatrix} \begin{pmatrix} 0 \\ \vdots \\ 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix} = \lambda \begin{pmatrix} 0 \\ \vdots \\ 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix} = \begin{pmatrix} 0 \\ \vdots \\ 0 \\ \lambda_k \\ 0 \\ \vdots \\ 0 \end{pmatrix} \quad (3.137)$$

woraus $\lambda = \lambda_k$ folgt. Die Eigenvektoren von M definieren aber die Orientierung der zu M korrespondierenden Ellipsoide, und da die Eigenvektoren offenbar durch

Abbildung 8: Links: Punktekonfiguration für $r_{xy} = .7$ mit Regressionsgeraden, Ellipsen und deren Hauptachsen; rechts: Die Hauptachsen als neue Koordinaten für die Punktekonfiguration.



die $\mathbf{e}_k$ gegeben sind, entsprechen sie den Orientierungen des Koordinatensystems, in dem die Ellipsoide liegen. Diese müssen also achsenparallel sein. $\square$

Definition 3.14 Die orthonormale Transformationsmatrix T rotiert das achsenparallele Ellipsoid $\mathcal{E}_y$ in das orientierte Ellipsoid $\mathcal{E}_x$, und wegen $T^{-1} = T'$ rotiert das Ellipsoid $\mathcal{E}_x$ in das Ellipsoid $\mathcal{E}_y$. T und T' heißen deshalb Hauptachsentransformationen.

Der mit der ersten Hauptachse der achsenparallelen Ellipse zusammenfallende Vektor $\mathbf{y}_1$ hat eine Länge, die sich aus der Ellipsengleichung

$$ay_1^2 + cy_2^2 = k > 0, \quad a > 0, \quad b > 0$$

ergibt (b ist ja für diese Ellipse gleich Null). Es ist aber $\mathbf{y}_1 = (y_1, 0)'$, so dass man insbesondere

$$\|\mathbf{y}_1\| = y_1 = \sqrt{k/a} \quad (3.138)$$

erhält. Analog dazu gilt

$$\|\mathbf{y}_2\| = y_2 = \sqrt{k/c}. \quad (3.139)$$

Offenbar sind die Längen der Hauptachsen umgekehrt proportional zu den Wurzeln aus den Eigenwerten. Aus (3.138) und (3.139) folgt

$$y_1 > y_2 \Rightarrow a < c, \quad y_1 < y_2 \Rightarrow a > c.$$

a und c , d.h. die Eigenwerte von M , spielen hier die Rolle von Skalenfaktoren. $a < c$ bedeutet, dass es eines größeren x_1 -Wertes bedarf, um eine vorgegebenen

Strecke auf der X_1 -Achse zu überdecken, als es einen x_2 -Wert benötigt, um die gleiche Strecke auf der X_2 -Achse zu überdecken.

Die Betrachtung überträgt sich ohne Weiteres auf den n -dimensionalen Fall, bei dem die Ellipse durch ein n -dimensionales Ellipsoid ersetzt wird.

Bemerkungen:

1. Da die Matrix T der Eigenvektoren einer symmetrischen Matrix M stets orthonormal ist, kann sie als eine Rotationsmatrix betrachtet werden, die die Vektoren $\mathbf{y} \in \mathcal{E}_y$ in die Vektoren $\mathbf{x} \in \mathcal{E}_x$ rotiert, wobei Λ die Diagonalmatrix der Eigenwerte von M ist. Umgekehrt rotiert T' die $\mathbf{x} \in \mathcal{E}_x$ in die Vektoren $\mathbf{y} \in \mathcal{E}_y$.
2. Eine Matrix muß nicht symmetrisch sein, damit Eigenvektoren für sie existieren. Allerdings existieren nicht für jede Matrix Eigenvektoren. Dazu betrachte man die Matrix (3.111), d.h. es sei

$$A = \begin{pmatrix} \cos \phi & -\sin \phi \\ \sin \phi & \cos \phi \end{pmatrix}.$$

Dann ist

$$A\mathbf{x} = x_1 \begin{pmatrix} \cos \phi \\ \sin \phi \end{pmatrix} + x_2 \begin{pmatrix} -\sin \phi \\ \cos \phi \end{pmatrix} = \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \mathbf{y},$$

und $\mathbf{y}$ ist nur parallel zu $\mathbf{x}$ für diejenigen Werte von ϕ , für die $\cos \phi = 1$ und $\sin \phi = 0$ ist, also z.B. für $\phi = 0$, so dass $A = I$ mit den Spaltenvektoren $(1, 0)'$ und $(0, 1)$. Dies ist der gewissermaßen triviale Fall, bei dem gar keine Rotation erzeugt wird. Man findet allerdings komplexwertige Eigenvektoren mit zugehörigen komplexwertigen Eigenwerten, – für $\phi = \pi/4$ etwa findet man die Eigenvektoren $(i, 1)'$ und $(-i, 1)'$ mit den Eigenwerten $(1 + i)/\sqrt{2}$ und $(1 - i)/\sqrt{2}$, mit $i = \sqrt{-1}$, wie man durch Nachrechnen bestätigt (s. s. Abschnitt 5.1, Seite 178). $\square$

3.9.3 Das charakteristische Polynom und Eigenräume

Es ist bisher nichts über die Möglichkeit, Eigenwerte und Eigenräume tatsächlich zu berechnen, gesagt worden, und insbesondere die Frage, ob alle Eigenwerte einer Matrix auch verschieden voneinander sind, ist nicht angesprochen worden. Die tatsächliche Berechnung geschieht mittels iterativer Verfahren, auf die hier nicht weiter eingegangen wird. Die zweite Frage läßt sich u. a. anhand der folgenden Betrachtungen angehen.

Es sei A eine $(n \times n)$ -Matrix und I die $(n \times n)$ -Einheitsmatrix, und es gelte

$$A\mathbf{v} = \lambda\mathbf{v}, \lambda \in \mathbb{R}, \mathbf{v} \in \mathbb{R}^n$$

d.h. λ ist ein Eigenwert von A und $\mathbf{v}$ ist der zugehörige Eigenvektor. Dann folgt

$$A\mathbf{v} - \lambda\mathbf{v} = (A - \lambda I)\mathbf{v} = \vec{0}. \quad (3.140)$$

Die rechte Seite zeigt ein homogenes lineares Gleichungssystem mit der Koeffizientenmatrix

$$\tilde{A} = \begin{pmatrix} a_{11} - \lambda_1 & a_{12} & a_{13} & \cdots & a_{1n} \\ a_{21} & a_{22} - \lambda_2 & a_{23} & \cdots & a_{2n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & a_{n3} & \cdots & a_{nn} - \lambda_n \end{pmatrix} \quad (3.141)$$

und dem Vektor $\mathbf{v}$ als Vektor von Unbekannten, das man aber nicht so ohne Weiteres lösen kann, weil λ ja ebenfalls nicht bekannt ist. Andererseits ist aus der Theorie der linearen Gleichungssysteme bekannt, dass $A\mathbf{x} = \vec{0}$ nur die Lösung $\mathbf{x} = \vec{0}$ hat, wenn A den vollen Rang hat; in dem Fall ist die Determinante ungleich Null. Damit also das System $(A - \lambda I)\mathbf{v} = \vec{0}$ mindestens eine von Null verschiedene Lösung hat, muß der Rang von $\tilde{A} = (A - \lambda I)$ kleiner als n sein, denn hätte $A - \lambda I$ den vollen Rang n , so kann der Nullvektor $\vec{0}$ in (3.140) als Linearkombination der Spalten von $A - \lambda I$ nur dargestellt werden, wenn $\mathbf{v} = \vec{0}$. Die Forderung $\mathbf{v} \neq \vec{0}$ impliziert also, dass $A - \lambda I$ nicht den vollen Rang hat, und in diesem Fall verschwindet die Determinante, so dass man auf die Gleichung

$$|A - \lambda I| = 0 \quad (3.142)$$

geführt wird. Entwickelt man diese Determinante, so ergibt sich ein Polynom in λ :

$$p_A = |A - \lambda I| = a_n \lambda^n + a_{n-1} \lambda^{n-1} + \cdots + a_1 \lambda + a_0 = 0. \quad (3.143)$$

Dies ist das *charakteristische Polynom* für die Eigenwertgleichung $A\mathbf{v} = \lambda\mathbf{v}$.

Wie sich zeigen läßt kann p_A in die Form

$$p_A(x) = \prod_{j=1}^n (x - \lambda_j) \quad (3.144)$$

gebracht werden. Eine Lösung λ_j bedeutet dann, dass für $x = \lambda_j$ das Polynom den Wert 0 annimmt, so dass die λ_j auch die Nullstellen des charakteristischen Polynoms heißen. Dem *Fundamentalsatz der Algebra* zufolge hat ein solches Polynom maximal n verschiedene Lösungen für λ , wobei für nichtsymmetrische Matrizen A auch komplexe Zahlen als Lösungen auftreten können. Die Eigenwerte λ_j können auch mit einer bestimmten Vielfachheit n_j auftreten; stets ist $\sum_j n_j = n$.

Ist λ_j eine Nullstelle, so kann man das Gleichungssystem

$$(A - \lambda_j I)\mathbf{v}_j = \vec{0} \quad (3.145)$$

lösen. $A - \lambda_j I$ ist eine Matrix und definiert damit eine Abbildung f von einem n -dimensionalen Vektorraum in einen n -dimensionalen Vektorraum, d.h. einen

Endomorphismus. Es werde angenommen, dass f injektiv sei; dann hat $A - \lambda_j I$ vollen Rang und (3.145) läßt nur den Nullvektor $\vec{0}$ als Lösung zu, und da $\vec{0}$ kein Eigenvektor ist erhalte man keinen Eigenvektor als Lösung. Mithin darf f *nicht* injektiv sein, damit die Lösung $\mathbf{v}_j$ ein Eigenvektor ist. Dieser Schluß folgt auch sofort aus dem Begriff der Injektivität: die Matrix $A - \lambda_j I$ bildet jeden Eigenvektor $\mathbf{v}_j$ auf den Nullvektor $\vec{0}$ ab, d.h. $f(\mathbf{v}_j) = \vec{0}$ für alle j . Die für die Injektivität charakteristische Relation $f(\mathbf{v}_j) = f(\mathbf{v}_k)$ für $j \neq k$ impliziert *nicht*, dass $\mathbf{v}_j = \mathbf{v}_k$.

Eine Lösung $\mathbf{v}_j$ definiert einen *Eigenraum*; für den Fall mehrfacher Lösungen $\mathbf{v}_j$ enthält der Eigenraum zu einem Eigenwert also mehr als einen Eigenvektor, – kommt ein Eigenvektor $\mathbf{v}_j$ mit der Vielfachheit n_j vor, so enthält der zugehörige Eigenraum n_j linear unabhängige zugehörige Eigenvektoren.

Beispiel 3.12 Eine (2×2) -Matrix Für eine (2×2) -Matrix mit den Diagonalelementen a und c ist oben gefunden worden, dass diese Diagonalelemente einerseits mit den Eigenwerten, andererseits mit den Längen der Hauptachsen ($\sqrt{k/a}$) und $\sqrt{k/c}$) in Beziehung stehen. Für $a = c$ sind die Eigenwerte gleich groß, und damit sind die Längen der Hauptachsen gleich groß, d.h. das Ellipsoid ist im Falle identischer Eigenwerte eine Kugel, was nahelegt, dass auch im Falle identischer Eigenwerte die zugehörigen Eigenvektoren orthogonal sind. Aus $MT = \Lambda T$ folgt dann wegen der Orthonormalität von T die Beziehung $MTT' = M = \Lambda$, d.h. M ist demnach eine Diagonalmatrix. Ist also $M = R$ eine Korrelationsmatrix, so ist diese im Falle identischer Eigenwerte eine Diagonalmatrix, d.h. die Variablen sind perfekt unkorreliert. Dann muß $r_{ii} = 1$ für alle i gelten und es folgt, dass die Eigenwerte alle gleich 1 sind. Es sind demnach die von Null verschiedenen Korrelationen, die ungleiche Eigenwerte implizieren. In Beispiel 3.9, Seite 103, werden die Eigenwerte einer (2×2) -Korrelationsmatrix R gegeben: sie haben die Werte $\lambda_1 = 1 + r$ und $\lambda_2 = 1 - r$. Offenbar $\lambda_1 = \lambda_2$ genau dann, wenn $r = 0$, – in diesem Fall ist die Korrelationsmatrix gleich der Identitätsmatrix. Es liegt nahe, ein solches Ergebnis für den $(n \times n)$ -Fall zu verallgemeinern: je ähnlicher die betrachtete symmetrische Matrix der Identitätsmatrix ist, desto ähnlicher werden die Eigenwerte. im Extremfall identischer Eigenwerte ist R eine Identitätsmatrix, und umgekehrt. $\square$

3.9.4 Spektraldarstellung einer symmetrischen Matrix M

Es sei M eine symmetrische $(n \times n)$ -Matrix mit der zugehörigen Matrix T von Eigenvektoren; Λ sei die Diagonalmatrix der zugehörigen Eigenwerte. Wegen der Orthonormalität der Vektoren in T folgt aus $MT = T\Lambda$ durch Multiplikation von links mit T'

$$T'MT = \Lambda. \quad (3.146)$$

Offenbar wird M durch Multiplikation von links mit T' und von rechts mit T in eine Diagonalmatrix, nämlich Λ , überführt, d.h. M wird *diagonalisiert*. Anderer-

seits folgt aus $MT = T\Lambda$ durch Multiplikation von rechts mit T'

$$M = T\Lambda T'. \quad (3.147)$$

Offenbar gilt für das Element m_{ij} von M dann die Beziehung

$$m_{ij} = \sum_{k=1}^n \lambda_k t_{ik} t_{jk}, \quad (3.148)$$

t_{ik}, t_{jk} die i -te Komponente des Eigenvektors $\mathbf{t}_k$, und t_{jk} die j -te Komponente von $\mathbf{t}_k$. Das Produkt $t_{ik} t_{jk}$ ist dann das Element in der i -ten Zeile und j -ten Spalte von $\mathbf{t}_k \mathbf{t}_k'$, so dass die Matrix M in der Form

$$M = \sum_{k=1}^n \lambda_k \mathbf{t}_k \mathbf{t}_k' \quad (3.149)$$

dargestellt werden kann; $\mathbf{t}_k \mathbf{t}_k'$ ist das dyadische Produkt des Vektors $\mathbf{t}_k$ mit sich selbst. Diese Darstellung heißt auch *Spektraldarstellung* von M .

3.9.5 Kovarianz und generalisierte Varianz

Es sei X eine zentrierte $(m \times n)$ -Datenmatrix, d.h. die Elemente von X seien $x_{ij} = X_{ij} - \bar{x}_j$; $X(i, j)$ ist das (i, j) -te Element der Matrix der Rohdaten. Die m Zeilen von X stehen für "Fälle" (Objekte, Personen), und die n Spalten von X repräsentieren Variablen. Dann ist

$$C = \frac{1}{m} X' X \quad (3.150)$$

die Varianz-Kovarianzmatrix für die Variablen. Sind die Elemente von X standardisiert worden, so gilt $x_{ij} = (X_{ij} - \bar{x}_j)/s_j$, wobei s_j die Standardabweichung der Messwerte in der j -ten Spalte der Datenmatrix ist, und

$$R = \frac{1}{m} X' X$$

eine Korrelationsmatrix. S und R sind symmetrisch.

Varianz-Kovarianz-Matrizen werden oft als Summe dyadischer Produkte dargestellt. Es sei $\tilde{\mathbf{x}}_i$ der i -te *Zeilenvektor* der Datenmatrix X , deren Spalten Variablen repräsentieren und deren Zeilen Individuen, Objekte, oder Zeitpunkte repräsentieren. Die Elemente x_{ij} der Matrix X seien zentriert. Die Kovarianz zwischen j -ter und k -ter Variable ist $s_{jk} = \sum_{i=1}^m x_{ij} x_{ik}$. Nun ist aber $x_{ij} x_{ik}$ das Element in der j -ten Zeile und k -ten Spalte der durch das dyadische Produkt $\tilde{\mathbf{x}}_i \tilde{\mathbf{x}}_i'$ definierten Matrix, so dass man auch

$$C = \sum_{i=1}^m \tilde{\mathbf{x}}_i \tilde{\mathbf{x}}_i' = \sum_{i=1}^m (\tilde{\mathbf{X}}_i - \bar{\mathbf{x}})(\tilde{\mathbf{X}}_i - \bar{\mathbf{x}})'. \quad (3.151)$$

Diese Gleichung bezieht sich auf eine gegebene Stichprobe von Messungen von n Variablen. Die Varianz-Kovarianz-Matrix wird oft auch anhand von Zufallsvektoren und ihrer Erwartungswerte definiert. So sei $\mathbf{X}' = (X_1, X_2, \dots, X_n)$ ein Zufallsvektor, dessen n Komponenten zufällige Variable sind, die die untersuchten (gemessenen) Variablen repräsentieren. Die Matrix X enthält in ihren Spalten die Realisierungen der Komponenten; x_{ij} ist die i -te Realisierung, also Messung, der j -ten Variablen. Der Vektor der Erwartungswerte ist

$$\mathbb{E}(\mathbf{X}) = \begin{pmatrix} \mathbb{E}(X_1) \\ \mathbb{E}(X_2) \\ \vdots \\ \mathbb{E}(X_n) \end{pmatrix} = \begin{pmatrix} \mu_1 \\ \mu_2 \\ \vdots \\ \mu_n \end{pmatrix} = \boldsymbol{\mu}. \quad (3.152)$$

Die Varianz-Kovarianz-Matrix wird dann als dyadisches Produkt von $\mathbf{X}$ mit sich selbst angeschrieben:

$$\boldsymbol{\Sigma} = \mathbb{E}[(\mathbf{X} - \boldsymbol{\mu})(\mathbf{X} - \boldsymbol{\mu})'] \quad (3.153)$$

Um klar zu machen, was hier gemeint ist, soll $\boldsymbol{\Sigma}$ ausgeschrieben werden:

$$\begin{aligned} \boldsymbol{\Sigma} &= \mathbb{E} \left[\begin{pmatrix} X_1 - \mu_1 \\ X_2 - \mu_2 \\ \vdots \\ X_n - \mu_n \end{pmatrix} (X_1 - \mu_1, X_2 - \mu_2, \dots, X_n - \mu_n) \right] \\ &= \mathbb{E} \begin{pmatrix} (X_1 - \mu_1)^2 & (X_1 - \mu_1)(X_2 - \mu_2) & \cdots & (X_1 - \mu_1)(X_n - \mu_n) \\ (X_2 - \mu_2)(X_1 - \mu_1) & (X_2 - \mu_2)^2 & \cdots & (X_2 - \mu_2)(X_n - \mu_n) \\ \vdots & \vdots & \ddots & \vdots \\ (X_n - \mu_n)(X_1 - \mu_1) & (X_n - \mu_n)(X_2 - \mu_2) & \cdots & (X_n - \mu_n)^2 \end{pmatrix} \\ &= \begin{pmatrix} \mathbb{E}(X_1 - \mu_1)^2 & \mathbb{E}(X_1 - \mu_1)(X_2 - \mu_2) & \cdots & \mathbb{E}(X_1 - \mu_1)(X_n - \mu_n) \\ \mathbb{E}(X_2 - \mu_2)(X_1 - \mu_1) & \mathbb{E}(X_2 - \mu_2)^2 & \cdots & \mathbb{E}(X_2 - \mu_2)(X_n - \mu_n) \\ \vdots & \vdots & \ddots & \vdots \\ \mathbb{E}(X_n - \mu_n)(X_1 - \mu_1) & \mathbb{E}(X_n - \mu_n)(X_2 - \mu_2) & \cdots & \mathbb{E}(X_n - \mu_n)^2 \end{pmatrix} \end{aligned}$$

Führt man die Bezeichnung $\sigma_{ij} = \mathbb{E}(X_i - \mu_i)(X_j - \mu_j)$ ein, so erhält man

$$\boldsymbol{\Sigma} = \begin{pmatrix} \sigma_{11} & \sigma_{12} & \cdots & \sigma_{1n} \\ \sigma_{21} & \sigma_{22} & \cdots & \sigma_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_{n1} & \sigma_{n2} & \cdots & \sigma_{nn} \end{pmatrix}. \quad (3.154)$$

Man kann C , wie in (3.151) erklärt, als Schätzung für $\boldsymbol{\Sigma}$ ansehen, – aber (3.150) ist dieselbe Schätzung, nur in anderer Schreibweise.

Der in der folgenden Definition eingeführte Begriff der verallgemeinerten Varianz ist gelegentlich nützlich

Definition 3.15 Die Determinante $|\Sigma|$ der Varianz-Kovarianz-Matrix Σ heißt generalisierte oder verallgemeinerte Varianz. Eine Schätzung der generalisierten Varianz ist $\frac{1}{n-1}|C|$, wobei C in (3.151) erklärt wurde.

Die geometrische Interpretation einer Determinante ist die eines Volumens. Die Menge der Vektoren $\mathbf{x}$, für die

$$(\mathbf{x} - \bar{\mathbf{x}})'C^{-1}(\mathbf{x} - \bar{\mathbf{x}}) = c^2 \quad (3.155)$$

gilt, definiert ein Ellipsoid (auch: Hyperellipsoid). Es kann gezeigt werden, dass

$$\text{Volumen von } \{\mathbf{x} | (\mathbf{x} - \bar{\mathbf{x}})'C^{-1}(\mathbf{x} - \bar{\mathbf{x}}) \leq c^2\} = \frac{2\pi^{n/2}|C|^{1/2}c^n}{n\Gamma(n/2)}, \quad (3.156)$$

mit

$$\Gamma(n/2) = \int_0^\infty t^{\frac{n-2}{2}} e^{-t} dt,$$

d.h. das Volumen ist proportional zur Wurzel aus der generalisierten Varianz. Die Bedeutung dieser Beziehung ist analog zu der einer einfachen Varianz σ^2 bzw. der Schätzung s^2 von σ^2 : die generalisierte Varianz gibt einen Eindruck von der Variabilität der Repräsentation der Variablen im n -dimensionalen Raum.

Soweit die Erläuterung einer Schreibweise für die Kovarianzmatrix. Im zweiten Teil dieses Abschnitts soll ein alternativer Beweis des Satzes ??, Seite ??, dass Kovarianzmatrizen positiv semidefinit sind, geliefert werden. Dieser soll der weiteren Einübung des Begriffs des Eigenvektors und dem mit ihm assoziierten Eigenwert dienen.

3.9.6 Die Inverse einer symmetrischen Matrix

Es sei $M \in \mathbb{R}^{n,n}$ eine symmetrische Matrix; M habe den Rang n , so dass M n Eigenwerte ungleich Null hat. Gesucht wird die zu M inverse Matrix M^{-1} . Sie kann leicht aus (3.147) bestimmt werden:

$$M^{-1} = (T\Lambda T')^{-1} = (T')^{-1}\Lambda^{-1}T^{-1}.$$

Wie oben bereits gezeigt ist aber $T^{-1} = T'$, und ebenso folgt $(T')^{-1} = (T^{-1})^{-1} = T$. Λ^{-1} enthält in den Diagonalelementen die Reziprokwerte der λ_j , als $1/\lambda_j$ (genau dann ist $\Lambda^{-1}\Lambda = I$), also folgt

$$M^{-1} = T\Lambda^{-1}T' = \sum_{k=1}^n \frac{\mathbf{t}_k \mathbf{t}_k'}{\lambda_k} \quad (3.157)$$

Anmerkung: Ein wichtiger Aspekt von (3.157) ist, dass *alle* Eigenwerte von M in die Definition der zu M inversen Matrix M^{-1} eingehen. Hat M nicht den Rang n , so ist mindestens einer der Eigenwerte gleich Null und (3.157) kann offenbar nicht

mehr berechnet werden, d.h. die inverse Matrix existiert dann nicht. Wenn die Matrix M vollen Rang hat, aber zumindest einige Eigenwerte klein sind, so werden die Elemente von M^{-1} groß. In der Regressionsstatistik lässt sich z.B. zeigen, dass die Varianzen der Schätzungen für Regressionskoeffizienten und die Kovarianz zwischen den Schätzungen für verschiedene Parameter proportional zu den Elementen der Inversen der Kovarianzmatrix für die ursprünglichen Messungen, also der Daten, sind. Sind die Kovarianzen zwischen den Messungen der Variablen hoch, werden zumindest einige der Eigenwerte der Varianz-Kovarianz-Matrix M klein. (3.157) impliziert, dass dann die Schätzungen der Regressionskoeffizienten mit großen Varianzen assoziiert sind, d.h. die Schätzungen sind dann instabil. $\square$

3.9.7 Die Wurzel aus einer positiv semidefiniten Matrix

Dass man Matrizen addieren und miteinander multiplizieren kann, ist bereits geklärt worden. Dass man auch eine Wurzel aus einer Matrix ziehen kann, zumindest, wenn sie symmetrisch und positiv semidefinit ist, soll noch kurz erwähnt werden. Es sei M eine symmetrische und positiv semidefinite $(n \times n)$ -Matrix; dann gilt

$$M = P\Lambda P', \quad \lambda_j \geq 0, \quad j = 1, \dots, n$$

und natürlich $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_n)$. Für Λ lässt sich die Wurzel $\Lambda^{1/2}$ leicht definieren. Es ist

$$\Lambda^{1/2} = \begin{pmatrix} \sqrt{\lambda_1} & 0 & \cdots & 0 \\ 0 & \sqrt{\lambda_2} & \cdots & 0 \\ & & \ddots & \\ 0 & 0 & \cdots & \sqrt{\lambda_n} \end{pmatrix}. \quad (3.158)$$

Sicherlich gilt dann $\Lambda = \Lambda^{1/2}\Lambda^{1/2}$. Man hat dann die

Definition 3.16 *Es sei M eine symmetrische und positiv semidefinite $(n \times n)$ -Matrix. Die Matrix*

$$M^{1/2} = P\Lambda^{1/2}P' = \sum_{k=1}^n \sqrt{\lambda_k} \mathbf{p}_k \mathbf{p}_k' \quad (3.159)$$

heißt Wurzel der Matrix M , wobei die $\mathbf{p}_k$ die Eigenvektoren von M sind und λ_k die zugehörigen Eigenwerte.

Man überzeugt sich sofort, dass diese Definition sinnvoll ist, denn es ist

$$M^{1/2}M^{1/2} = P\Lambda^{1/2}P\Lambda^{1/2}P' = P\Lambda P' = M.$$

Es folgt sofort aus (3.159), dass $M^{1/2}$ symmetrisch ist, denn

$$(M^{1/2})' = (P\Lambda^{1/2}P')' = P\Lambda^{1/2}P'.$$

3.9.8 Die Singularwertzerlegung (SVD)

Es sei $X = [\mathbf{x}_1, \dots, \mathbf{x}_n]$ eine $(m \times n)$ -Matrix $X \in \mathbb{R}^{m,n}$ mit dem Rang $r \leq \min(m, n)$. Es werden r m -dimensionale Basisvektoren $\mathbf{L}_1, \dots, \mathbf{L}_r$ gesucht derart, dass die m -dimensionalen Spaltenvektoren $\mathbf{x}_j$ von X als Linearkombinationen der $\mathbf{L}_k$ dargestellt werden können. Dies führt auf den Ansatz

$$X = LA, \quad (3.160)$$

A eine Matrix von Koeffizienten. Ist $\mathbf{a}_j$ der j -te Spaltenvektor von A , so gilt

$$\mathbf{x}_j = L\mathbf{a}_j, \quad j = 1, \dots, n \quad (3.161)$$

Auf den ersten Blick mag es als notwendig erscheinen, zwei Matrizen – L und A – bestimmen zu müssen. Nimmt man allerdings an, dass A derart definiert werden kann, dass die Inverse A^{-1} existiert, so liefert (3.160) nach Multiplikation von rechts mit $B = A^{-1}$ die Beziehung

$$XB = L. \quad (3.162)$$

Hat man also eine Matrix B mit existierender Inverse $B^{-1} = A$ gewählt, so ergeben sich die Spalten $\mathbf{L}_k$ von L als Linearkombinationen der Spalten $\mathbf{x}_j$ von X :

$$\mathbf{L}_k = X\mathbf{b}_k, \quad k = 1, \dots, r \quad (3.163)$$

$\mathbf{b}_k$ die k -te Spalte von B . Die Frage ist nun, nach welchen Kriterien die Matrix B gewählt werden soll.

Es ist an dieser Stelle sinnvoll, sich die Eigenvektorgleichungen für die Kreuzproduktmatrizen $X'X$ und XX' zu vergegenwärtigen, wobei $\text{rg}(X) = \text{rg}(X'X) = \text{rg}(XX') \leq \min(m, n)$ gelte (vergl. (3.65), Seite 80):

$$(X'X)T_r = T_r\Lambda_{r1}, \quad (3.164)$$

$$(XX')Q_r = Q_r\Lambda_{r2}. \quad (3.165)$$

T_r und Q_r enthalten die zu Eigenwerten ungleich Null korrespondierenden Eigenvektoren. T_r ist eine $(n \times r)$ -Matrix und Q_r ist eine $(m \times r)$ -Matrix. Λ_{r1} und Λ_{r2} enthalten ebenfalls nur die von Null verschiedenen r Eigenwerte. Dann gilt der

Satz 3.20 *Es sei entsprechend (3.162) mit $B = T_r$*

$$XT_r = L_r. \quad (3.166)$$

Die von Null verschiedenen Eigenwerte von $X'X$ und XX' sind identisch, d.h. es gilt $\Lambda_{r1} = \Lambda_{r2}$, und die normierten Spaltenvektoren von $S_r = L_r\Lambda_{r1}^{-1/2}$ sind die Eigenvektoren von XX' , d.h. es gilt $S_r = Q_r$.

Beweis: Die Matrix L_r ist sicher orthogonal, denn

$$L_r' L_r = T_r' X' X T_r = T_r' T_r \Lambda_{r1} T_r' T_r = \Lambda_{r1} = \text{diag}(\lambda_1, \dots, \lambda_r).$$

Dieses Ergebnis bedeutet

$$\mathbf{L}_k' \mathbf{L}_k = \|\mathbf{L}_k\|^2 = \lambda_k, \quad k = 1, \dots, r \quad (3.167)$$

d.h. für die Länge von $\mathbf{L}_k$ gilt $\|\mathbf{L}_k\| = \sqrt{\lambda_k}$. $\mathbf{L}_k$ kann normiert werden, indem $\mathbf{L}_k$ mit $1/\sqrt{\lambda_k}$ multipliziert wird:

$$\mathbf{s}_k = \frac{1}{\sqrt{\lambda_k}} \mathbf{L}_k, \text{ bzw. } S_r = L_r \Lambda_{r1}^{-1/2}. \quad (3.168)$$

S_r ist sicher orthonormal, denn

$$S_r' S_r = \Lambda_{r1}^{-1/2} L_r' L_r \Lambda_{r1}^{-1/2} = \Lambda_{r1}^{-1/2} \Lambda_{r1} \Lambda_{r1}^{-1/2} = \Lambda_{r1}^{-1/2} \Lambda_{r1}^{1/2} \Lambda_{r1}^{1/2} \Lambda_{r1}^{-1/2} = I_r,$$

I_r die $(r \times r)$ -Einheitsmatrix. Dann folgt $L_r = S_r \Lambda_{r1}^{1/2}$ und X läßt sich in der Form

$$X = S_r \Lambda_{r1}^{1/2} T_r' \quad (3.169)$$

darstellen. Für XX' erhält man den Ausdruck

$$XX' = S_r \Lambda_{r1}^{1/2} T_r' T_r \Lambda_{r1}^{1/2} S_r' = S_r \Lambda_{r1} S_r'. \quad (3.170)$$

Wegen der Orthonormalität von S_r folgt nach Multiplikation dieser Gleichung von rechts mit S_r

$$(XX') S_r = S_r \Lambda_{r1}, \quad (3.171)$$

d.h. S_r enthält die zu Eigenwerten ungleich Null korrespondierenden Eigenvektoren von XX' . Der Vergleich mit (3.165) zeigt, dass $Q_r = S_r$ und $\Lambda_{r2} = \Lambda_{r1}$.

□

Die Gleichung (3.169) ist als *Singularwertzerlegung* (SVD) der Matrix X bekannt, allerdings wird sie üblicherweise in etwas anderer Form angeschrieben. Es sei Q die $(m \times m)$ -Matrix der Matrix XX' , d.h. Q enthält auch die $n - r$ Eigenvektoren von XX' , die zu den Eigenwerten gleich Null von XX' korrespondieren, falls $r < \min(m, n)$. Analog dazu sei T die Matrix aller Eigenvektoren von $X'X$, also einschließlich der Eigenvektoren, die zu Eigenwerten gleich Null von $X'X$ korrespondieren, falls $r < \min(m, n)$. Ferner sei

$$\Sigma = \Lambda^{1/2} = \begin{pmatrix} \Lambda_r^{1/2} & 0 \\ 0 & 0 \end{pmatrix}, \quad (3.172)$$

wobei Λ_r die von Null verschiedenen Eigenwerte von $X'X$ bzw. XX' enthält und die Nullen die $n - r$ Zeilen bzw. Spalten von Nullen repräsentieren, falls $r < \min(m, n)$. Dann hat man die

Definition 3.17 *Die Darstellung*

$$X = Q\Sigma T' \quad (3.173)$$

heißt Singularwertzerlegung von X . Die Spaltenvektoren von Q heißen Linkssingularvektoren, die von T heißen Rechtssingularvektoren, und die Diagonalelemente $\sigma_j = \sqrt{\lambda_j}$ von Σ heißen Singularwerte.

Anmerkungen:

1. In Abschnitt 3.6 wurde für eine beliebige Matrix X mit dem Rang $\text{rg}(X)$ die Beziehung (3.68) (Seite 81), also $X = UV$ hergeleitet, wobei U eine Matrix mit r linear unabhängigen m -dimensionalen Vektoren, V eine $(r \times n)$ -Matrix mit r linear unabhängigen, n -dimensionalen Vektoren ist. Die SVD kann als Spezialfall dieser Gleichung mit entweder $U = Q$, $V = \Lambda^{1/2}T'$ oder $U = Q\Lambda^{1/2}$, $V = T'$ gesehen werden. $\square$
2. Nach Definition 3.14, Seite 107, repräsentiert die Matrix T der Eigenvektoren von $X'X$ eine Hauptachsentransformation der Spalten von X . Die Singularwertzerlegung von X ist äquivalent einer Hauptachsentransformation, denn $XT = Q\Lambda^{1/2} = L$.
3. Im Satz 3.2 (Seite 78) wurde ausgesagt, dass jede $(m \times n)$ -Matrix X als Produkt zweier Matrizen U und V' dargestellt werden kann, $X = UV'$, wobei $\text{rg}(X) = \text{rg}(U) = \text{rg}(V)$. Die SVD (3.173) ist ein solches Produkt; man kann entweder $U = Q$ und $V' = \Sigma T'$ bzw. $V = T\Sigma$ oder $U = Q\Sigma$ und $V = T$ setzen.
4. Die englische Bezeichnung für 'Singularwertzerlegung' ist *singular value decomposition*, abgekürzt SVD; diese Abkürzung ist auch im Deutschen üblich. Ein in der Psychologie häufig gebrauchter Ausdruck für die SVD ist 'Grundstruktur' einer Matrix (engl. basic structure). Der Ausdruck 'Singularwertzerlegung' ist allgemein in allen Wissenschaften, in denen eine Zerlegung von Matrizen gewünscht wird (Biologie, Medizin, Geologie, Klimaforschung, Archäologie, etc) gebräuchlich, weshalb auch hier von dieser Bezeichnung Gebrauch gemacht wird.
5. Die SVD ist nicht an eine Spaltenzentrierung oder Standardisierung der Matrix X gebunden; eine SVD kann für eine beliebige Matrix X bestimmt werden. Ist X eine $(m \times n)$ -Datenmatrix, bei der die Zeilen Fälle und die Spalten Variablen repräsentieren, so hat man im Allgemeinen $m > n$; dies ist keine Bedingung für die SVD von X . $\square$

Aus (3.173) folgen die Gleichungen

$$XT = Q\Sigma \quad (3.174)$$

$$Q'X = \Sigma T', \text{ d.h. } X'Q = T\Sigma'. \quad (3.175)$$

Diese beiden Gleichungen verdeutlichen die Beziehung zwischen den Eigenvektoren in T und Q . Es sei $\mathbf{q}_i$ der i -te Spaltenvektor von Q und $\mathbf{t}_i$ der i -te Spaltenvektor von T , die Gleichungen (3.174) und (3.175) implizieren dann

$$X\mathbf{t}_i = \sigma_i\mathbf{q}_i, \quad i = 1, \dots, m \quad (3.176)$$

$$X'\mathbf{q}_i = \sigma_i\mathbf{t}_i \quad (3.177)$$

woraus

$$X'X\mathbf{t}_i = \sigma_i X'\mathbf{q}_i = \sigma_i^2\mathbf{t}_i \quad (3.178)$$

$$XX'\mathbf{q}_i = \sigma_i X\mathbf{t}_i = \sigma_i^2\mathbf{q}_i. \quad (3.179)$$

folgt.

Darstellung der SVD über das dyadische Produkt: Die Zerlegung $X = Q\Sigma T'$ kann in der Form

$$X = \sum_{k=1}^n \sigma_k \mathbf{q}_k \mathbf{t}'_k = \sum_{k=1}^n \sqrt{\lambda_k} \mathbf{q}_k \mathbf{t}'_k, \quad m \geq n \quad (3.180)$$

dargestellt werden, wobei $\mathbf{q}_j$ die Spaltenvektoren von Q und $\mathbf{t}_j$ die Spaltenvektoren von T sind. $\mathbf{q}_j \mathbf{t}'_j$ ist das dyadische Produkt dieser Vektoren. X kann also als Summe von Matrizen aufgefasst werden, die jeweils eine Dimension repräsentieren. Das Element x_{ij} ist demnach durch die Summe

$$x_{ij} = \sqrt{\lambda_1} q_{i1} t_{j1} + \dots + \sqrt{\lambda_n} q_{in} t_{jn} \quad (3.181)$$

gegeben, $\sqrt{\lambda_1} q_{i1} t_{j1}$ ist der Beitrag der ersten latenten Dimension, etc.

Zentrierung: Es sei X eine Datenmatrix, bei der die Spalten Messwerte für bestimmte Variablen repräsentieren. Ist man in erster Linie an der Analyse der Beziehungen zwischen den Variablen – repräsentiert durch die Spalten von X – interessiert, wird man die Spalten von X zentrieren, d.h. man wird z.B. von den Messwerten X_{ij} zu den Abweichungen $x_{ij} = X_{ij} - \bar{x}_j$ übergehen, wobei $\bar{x}_j = \sum_i X_{ij}/m$ der Mittelwert der Werte in der j -ten Spalte ist. Die x_{ij} heißen *zentrierte* Werte. Dividiert man überdies die x_{ij} noch durch die Standardabweichungen s_j der Werte in der j -ten Spalten, erhält man *spaltenstandardisierte* Werte z_{ij} . Es sei $\vec{1}$ der Einsvektor, d.h. der m -dimensionale Vektor, dessen Komponenten alle gleich 1 sind. Dann ist

$$\bar{\mathbf{x}} = \frac{1}{m} X' \vec{1} \quad (3.182)$$

der Vektor der Spaltenmittelwerte. Enthält X bereits zentrierte Werte, so ist $\bar{\mathbf{x}} = \vec{0}$ der Nullvektor (die Summe der Abweichungen vom Mittelwert ist stets gleich Null). Insbesondere gilt $Z' \vec{1} = \vec{0}$. Für (spalten-)zentrierte oder (spalten-)standardisierte Werte folgt dann Für

$$T' \Sigma X' \vec{1} = Q' \vec{1} = \vec{0}, \quad (3.183)$$

$$T' X' \vec{1} = L' \vec{1} = \vec{0}, \quad (3.184)$$

Zur Interpretation der Eigenwerte: Aus dem Vorangegangenen folgt

$$\frac{1}{m} \|\mathbf{q}_k\|^2 = 1 \quad (3.185)$$

$$\frac{1}{m} \|\mathbf{L}_k\|^2 = s_k^2 \quad (3.186)$$

wobei s_k^2 die Varianz der Komponenten von $\mathbf{L}_k$ ist. Wegen

$$T'X'XT = L'L = \Lambda = \text{diag}(\|\mathbf{L}_1\|^2, \dots, \|\mathbf{L}_n\|^2) = \text{diag}(\lambda_1, \dots, \lambda_n) \quad (3.187)$$

sind deshalb die Eigenwerte λ_k von $X'X$ proportional zu den Varianzen s_k^2 der Koordinaten der Fälle auf den latenten Variablen, also

$$\frac{1}{m} \|\mathbf{L}_k\|^2 = s_k^2 = \frac{1}{m} \lambda_k \quad (3.188)$$

Analoge Betrachtungen gelten für den Fall der Zeilenzentrierung; man geht dann von den Messwerten X_{ij} zu den $x_{ij} = X_{ij} - \bar{x}_{(i)}$ über, wobei $\bar{x}_{(i)}$ der Mittelwert der i -ten Zeile ist. Dieser Ansatz setzt voraus, dass es Sinn macht, über die Variablen zu mitteln. Sind die Variablen verschiedene Größen wie (i) galvanischer Hautwiderstand, (ii) Hormonkonzentration im Speichel, (iii) Herzrate, etc., so ist nicht ganz klar, was $\bar{x}_{(i)}$ bedeuten soll.

3.10 Maximalprinzipien

3.10.1 Die Differentiation von Vektoren

Ein Vektor ist durch seine Komponenten festgelegt. Man kann dann fragen, wie sich der Vektor verändert, wenn man seine Komponenten verändert. Solche Veränderungen lassen sich oft durch einen Differentialquotienten beschreiben. So sei etwa $\mathbf{x} = (x_1, x_2, \dots, x_n)'$. Dabei wird stillschweigend angenommen, dass keine Komponente von der anderen abhängt. Man definiert nun den Differentialquotienten von $\mathbf{x}$ in Bezug auf die j -te Komponente x_j durch

$$\frac{d\mathbf{x}}{dx_j} = \begin{pmatrix} dx_1/dx_j \\ \vdots \\ dx_j/dx_j \\ \vdots \\ dx_n/dx_j \end{pmatrix} = \begin{pmatrix} 0 \\ \vdots \\ 1 \\ \vdots \\ 0 \end{pmatrix} = \mathbf{e}_j. \quad (3.189)$$

Es gibt noch einen zweiten Fall, bei dem die Komponente von einer Variablen, etwa der Zeit t , abhängen, so dass

$$\mathbf{x}(t) = \begin{pmatrix} x_1(t) \\ x_2(t) \\ \vdots \\ x_n(t) \end{pmatrix}$$

geschrieben wird. Man kann dann die Veränderung von $\mathbf{x}$ mit t durch den Vektor

$$\frac{d\mathbf{x}(t)}{dt} = \begin{pmatrix} dx_1(t)/dt \\ dx_2(t)/dt \\ \vdots \\ dx_n(t)/dt \end{pmatrix}$$

ausdrücken. Dieser Fall wird im Folgenden nicht behandelt.

Der Fall (3.189) tritt u.a. dann auf, wenn eine Größe in Abhängigkeit von einem Vektor $\mathbf{b} = (b_1, \dots, b_p)'$ von Parametern maximiert oder minimiert werden soll.

Es sei $\mathbf{y} = A\mathbf{x}$; für eine gegebene Matrix A hängt $\mathbf{y}$ von $\mathbf{x}$ ab, so dass man $\mathbf{y} = \mathbf{y}(\mathbf{x})$ schreiben kann. Man findet dann als unmittelbare Konsequenz aus (3.189)

$$\frac{\partial \mathbf{y}}{\partial \mathbf{x}} = A. \quad (3.190)$$

3.10.2 Die Differentiation von quadratischen Formen

Bei der Schätzung von Parametern kommt es immer wieder vor, quadratische Formen $Q(\mathbf{x}) = \mathbf{x}'C\mathbf{x}$, wobei C eine symmetrische $(n \times n)$ -Matrix ist, nach dem n -dimensionalen Vektor $\mathbf{x}$ zu differenzieren. Man differenziert zunächst nach der Komponente x_j von $\mathbf{x}$ und findet (Kettenregel)

$$\frac{\partial Q}{\partial x_j} = \mathbf{e}_j' C \mathbf{x} + \mathbf{x}' C \mathbf{e}_j. \quad (3.191)$$

hierin sind $\mathbf{e}_j' C \mathbf{x}$ und $\mathbf{x}' C \mathbf{e}_j$ Skalare. Nun ist $\mathbf{e}_j' C$ gleich dem j -ten Zeilenvektor $(c_{j1}, c_{j2}, \dots, c_{jn})$ von C , so dass

$$\mathbf{e}_j' C \mathbf{x} = \sum_{k=1}^n c_{jk} x_k = \sum_{k=1}^n x_k c_{jk}.$$

Weiter ist $C \mathbf{e}_j$ gleich der j -ten Spalte von C und

$$\mathbf{x}' C \mathbf{e}_j = \sum_{k=1}^n x_j c_{kj}.$$

Wegen der Symmetrie von C sind aber die j -te Zeile und die j -te Spalte von C identisch, so dass $\mathbf{e}_j' C \mathbf{x} = \mathbf{x}' C \mathbf{e}_j$ und

$$\frac{\partial Q}{\partial x_j} = 2\mathbf{e}_j' C \mathbf{x} = 2\mathbf{x}' C \mathbf{e}_j \quad (3.192)$$

Fasst man die Ausdrücke für alle j zusammen, so erhält man

$$\frac{\partial Q}{\partial \mathbf{x}} = 2C\mathbf{x} \quad \text{bzw.} \quad \frac{\partial Q}{\partial \mathbf{x}} = 2\mathbf{x}' C, \quad (3.193)$$

denn das Zusammenfassen der Einheitsvektoren $\mathbf{e}_j$ führt auf die Einheitsmatrix. Damit hat man auch die Ableitung von $\|\mathbf{x}\|^2 = \mathbf{x}'\mathbf{x}$ nach $\mathbf{x}$ gefunden: es ist ja $\mathbf{x}'\mathbf{x} = \mathbf{x}'C\mathbf{x}$ mit $C = I$, I die Identitätsmatrix. Dann folgt

$$\frac{\partial \mathbf{x}'\mathbf{x}}{\partial \mathbf{x}} = 2\mathbf{x}. \quad (3.194)$$

Beispiel 3.13 Es gelte $\mathbf{x}'M\mathbf{x} = k$ mit $M' = M$ und $\mathbf{y}'\Lambda\mathbf{y} = k$, Λ eine Diagonalmatrix. Weiter sei $\mathbf{x} = T\mathbf{y}$, $T'T = I$. Dann folgt

$$\mathbf{x}'M\mathbf{x} = \mathbf{y}'T'MT\mathbf{y} = \mathbf{y}'\Lambda\mathbf{y} = k, \quad \mathbf{y} \in \mathcal{E}_y,$$

$\mathcal{E}_y$ das durch Λ definierte Ellipsoid. Dann muß

$$\frac{d(\mathbf{y}'T'MT\mathbf{y})}{d\mathbf{y}} = \frac{d(\mathbf{y}'\Lambda\mathbf{y})}{d\mathbf{y}}$$

gelten. Nach (3.193) muß dann

$$2T'MT\mathbf{y} = 2\Lambda\mathbf{y}$$

gelten, und nochmalige Differentiation nach $\mathbf{y}$ liefert dann

$$T'MT = \Lambda,$$

(vergl. (3.190)); diese Aussage wurde bereits auf Seite 99, wurde die Aussage bereits ohne weitere Herleitung gefolgert (vergl. Gleichung (3.116)). $\square$

3.10.3 Die Methode der Kleinsten Quadrate

Es sei

$$\mathbf{y} = X\mathbf{b} + \mathbf{e}. \quad (3.195)$$

Es wird im Allgemeinen angenommen, dass die Varianz-Kovarianz-Matrix der Fehler $\mathbf{e}$ durch

$$\mathbb{V}(\mathbf{e}) = \sigma^2 I \quad (3.196)$$

gegeben ist, I die Einheitsmatrix, d.h. der Fehler ist für alle Komponenten von $\mathbf{y}$ gleich und die Kovarianzen der Fehlerkomponenten sind alle gleich Null. Es gibt zwei Möglichkeiten:

1. Die Spaltenvektoren $\mathbf{x}_1, \dots, \mathbf{x}_p$ von X enthalten m Messwerte von p Variablen. $X\mathbf{b}$ ist dann die Linearkombination

$$\mathbf{y}_0 = X\mathbf{b} = b_1\mathbf{x}_1 + b_2\mathbf{x}_2 + \dots + b_p\mathbf{x}_p. \quad (3.197)$$

2. Der erste Spaltenvektor von X ist $\mathbf{1}_m = (1, 1, \dots, 1)'$, d.h. die erste Spalte von X

enthält nur Einsen. $\mathbf{b}$ ist dann ein $(p+1)$ -dimensionaler Vektor $\mathbf{b} = (b_0, b_1, \dots, b_p)'$ und $\mathbf{y}_0$ ist durch

$$\mathbf{y}_0 = X\mathbf{b} = b_0 + b_1\mathbf{x}_1 + \dots + b_p\mathbf{x}_p \quad (3.198)$$

definiert.

X ist also eine $(m \times p)$ - oder $(m \times p+1)$ -dimensionale Matrix, wobei $p < m$ bzw. $p+1 < m$ vorausgesetzt wird, $\mathbf{y}$ und $\mathbf{e}$ sind m -dimensionale Vektoren und $\mathbf{b}$ ist ein p - bzw. $(p+1)$ -dimensionaler Vektor. Dies ist das Allgemeine Lineare Modell; $\mathbf{y}$ und X sind gegeben, $\mathbf{b}$ und $\mathbf{e}$ sind unbekannt, und es wird eine Schätzung $\hat{\mathbf{b}}$ für $\mathbf{b}$ gesucht derart, dass $\mathbf{y} = X\hat{\mathbf{b}} + \hat{\mathbf{e}}$ und $\hat{\mathbf{e}}'\hat{\mathbf{e}}$ minimal ist. $\hat{\mathbf{e}}$ ist im Allgemeinen nicht gleich $\mathbf{e}$, sondern hängt von $\hat{\mathbf{b}}$ ab. Nach der Methode der Kleinsten Quadrate wird $\hat{\mathbf{b}}$ so bestimmt, dass $\hat{\mathbf{e}}'\hat{\mathbf{e}}$ minimal wird.

Nach (3.195) ist

$$\|\mathbf{e}\|^2 = \mathbf{e}'\mathbf{e} = (\mathbf{y} - X\mathbf{b})'(\mathbf{y} - X\mathbf{b}) = Q(\mathbf{b}),$$

mit $\mathbf{e} = \mathbf{y} - X\mathbf{b}$, und $Q(\mathbf{b})$ nimmt einen extremen Wert an, wenn die Ableitung $dQ/d\mathbf{b}$ gleich Null ist. Nach der Kettenregel folgt

$$\frac{dQ}{d\mathbf{b}} = \frac{dQ}{d\mathbf{e}} \frac{d\mathbf{e}}{d\mathbf{b}} = -2\mathbf{e}'X, \quad (3.199)$$

nach (3.190) und (3.193), und

$$\left. \frac{dQ}{d\mathbf{b}} \right|_{\mathbf{b}=\hat{\mathbf{b}}} = (\mathbf{y} - X\hat{\mathbf{b}})'X = 0, \quad (3.200)$$

d.h. der Vektor $\mathbf{y} - X\hat{\mathbf{b}}$ ist orthogonal zu allen Spaltenvektoren von X , so dass

$$\mathbf{y}'X = \hat{\mathbf{b}}X'X,$$

woraus sofort

$$\hat{\mathbf{b}} = (X'X)^{-1}X'\mathbf{y} \quad (3.201)$$

als Kleinste-Quadrate-Schätzung für $\mathbf{b}$ folgt. Die stillschweigende Voraussetzung ist, dass die Inverse $(X'X)^{-1}$ existiert.

Anmerkungen: In Bezug auf (3.200) wurde angemerkt, dass $\mathbf{y} - X\hat{\mathbf{b}}$ orthogonal zu allen Spaltenvektoren von X ist. Es ist aber, mit $\hat{\mathbf{y}} = X\hat{\mathbf{b}}$,

$$\mathbf{y} - X\hat{\mathbf{b}} = \mathbf{y} - \hat{\mathbf{y}} = \hat{\mathbf{e}} \quad (3.202)$$

d.h. der zu $\hat{\mathbf{b}}$ korrespondierende Fehlervektor $\hat{\mathbf{e}}$ ist orthogonal zu jedem Spaltenvektor $\mathbf{x}_j$ von X . Darüber hinaus gilt

$$\hat{\mathbf{e}}'\hat{\mathbf{y}} = 0, \quad (3.203)$$

d.h. der Fehlervektor $\hat{\mathbf{e}}$ ist orthogonal zu $\hat{\mathbf{y}}$, wie man leicht sieht: Es ist wegen (3.201)

$$\hat{\mathbf{y}} = X\hat{\mathbf{b}} = X(X'X)^{-1}X'\mathbf{y}, \quad (3.204)$$

so dass

$$\begin{aligned}\hat{\mathbf{e}}'\hat{\mathbf{y}} &= (\mathbf{y} - \hat{\mathbf{y}})'\hat{\mathbf{y}} = \mathbf{y}'\hat{\mathbf{y}} - \hat{\mathbf{y}}'\hat{\mathbf{y}} \\ &= \mathbf{y}'X(X'X)^{-1}X'\mathbf{y} - \mathbf{y}'X(X'X)^{-1}X'X(X'X)^{-1}\mathbf{y} \\ &= \mathbf{y}'X(X'X)^{-1}X'\mathbf{y} - \mathbf{y}'X(X'X)^{-1}X'\mathbf{y} = 0.\end{aligned}\quad (3.205)$$

Es sei $P_r = X(X'X)^{-1}X'$. Offenbar gilt $P_r P_r = P_r$; P_r ist ein Beispiel für eine *idempotente* Matrix.

Zur Übung werde die einfache lineare Regression betrachtet, bei der $y = b_1x + b_0 + e$ mit b_1 und b_0 als unbekannten Konstanten gilt. Für eine gegebene Stichprobe von x - und y -Werten kann man dann

$$\mathbf{y} = b_0 + b_1\mathbf{x} + \mathbf{e} \quad (3.206)$$

schreiben. Setzt man $X = [\vec{1}, \mathbf{x}]$, so dass X eine Matrix ist, deren erste Spalte nur Einsen und deren zweite Spalte die gemessenen x -Werte enthält, so liefert (3.201) die Schätzungen

$$\hat{\mathbf{b}} = \begin{pmatrix} \hat{b}_0 \\ \hat{b}_1 \end{pmatrix} = \begin{pmatrix} m, & \sum_i x_i \\ \sum_i x_i, & \sum_i x_i^2 \end{pmatrix}^{-1} X'\mathbf{y} \quad (3.207)$$

Hierin ist

$$X'\mathbf{y} = \begin{pmatrix} \sum_i y_i \\ \sum_i x_i y_i \end{pmatrix}$$

und nach (3.95), Seite 91, ist

$$\begin{pmatrix} m, & \sum_i x_i \\ \sum_i x_i, & \sum_i x_i^2 \end{pmatrix}^{-1} = \frac{1}{m \sum_i x_i^2 - (\sum_i x_i)^2} \begin{pmatrix} \sum_i x_i^2, & -\sum_i x_i \\ -\sum_i x_i, & m \end{pmatrix} \quad (3.208)$$

Für den Faktor im Ausdruck für die Inverse erhält man nach Multiplikation von Zähler und Nenner mit $1/m^2$

$$\alpha = \frac{1}{m \sum_i x_i^2 - (\sum_i x_i)^2} = \frac{1/m^2}{\frac{1}{m} \sum_i x_i^2 - \bar{x}^2} = \frac{1/m^2}{s_x^2},$$

und

$$\hat{\mathbf{b}} = \alpha \begin{pmatrix} \sum_i x_i^2, & -\sum_i x_i \\ -\sum_i x_i, & m \end{pmatrix} \begin{pmatrix} \sum_i y_i \\ \sum_i x_i y_i \end{pmatrix} = \alpha \begin{pmatrix} \sum_i x_i^2 \sum_i y_i - \sum_i x_i \sum_i x_i y_i \\ -\sum_i x_i \sum_i y_i + m \sum_i x_i y_i \end{pmatrix}.$$

Diese Gleichung liefert

$$\hat{b}_0 = \frac{\bar{y} \frac{1}{m} \sum_i x_i^2 - \bar{x} \frac{1}{m} \sum_i x_i y_i}{s_x^2} \quad (3.209)$$

$$\begin{aligned}\hat{b}_1 &= \alpha (m \sum_i x_i y_i - \sum_i x_i \sum_i y_i) = \frac{\frac{1}{m} \sum_i x_i y_i - \bar{x} \bar{y}}{s_x^2} \\ &= \frac{Kov(x, y)}{s_x^2}.\end{aligned}\quad (3.210)$$

Der übliche Ausdruck für $\hat{b}_0$ in der einfachen Regression ist $\hat{b}_0 = \bar{y} - \hat{b}_1 \bar{x}$, und der Ausdruck in (3.209) scheint davon abzuweichen, aber die beiden Ausdrücke sind identisch. Um das zu sehen, muß man den Ausdruck (3.209) erweitern:

$$\begin{aligned}
\hat{b}_0 &= \frac{\bar{y} \frac{1}{m} \sum_i x_i^2}{s_x^2} - \frac{\bar{x} \frac{1}{m} \sum_i x_i y_i}{s_x^2} \\
&= \frac{\bar{y} \frac{1}{m} \sum_i x_i^2 - \bar{x}^2 \bar{y} + \bar{x}^2 \bar{y}}{s_x^2} - \frac{\bar{x} \frac{1}{m} \sum_i x_i y_i - \bar{x}^2 \bar{y} + \bar{x}^2 \bar{y}}{s_x^2}, \\
&= \frac{\bar{y} (\frac{1}{m} \sum_i x_i^2 - \bar{x}^2)}{s_x^2} + \frac{\bar{x}^2 \bar{y}}{s_x^2} - \frac{\bar{x} (\frac{1}{m} \sum_i x_i y_i - \bar{x} \bar{y})}{s_x^2} - \frac{\bar{x}^2 \bar{y}}{s_x^2} \\
&= \bar{y} - \frac{Kov(x, y)}{s_x^2} \bar{x} = \bar{y} - \hat{b}_1 \bar{x},
\end{aligned} \tag{3.211}$$

in Übereinstimmung mit dem üblichen Ausdruck, da ja $\frac{1}{m} \sum_i x_i^2 - \bar{x}^2 = s_x^2$.

3.10.4 Generalisierte Kleinste Quadrate

Statt (3.196) gelte nun

$$\mathbb{V}(\mathbf{e}) = \sigma^2 \Sigma, \tag{3.212}$$

wobei $\sigma^2 \in \mathbb{R}$ unbekannt sei, die Varianz-Kovarianz-Matrix Σ aber bekannt sei. Σ ist symmetrisch und werde als positiv-definit vorausgesetzt. Wenn, entgegen der Annahme, die Fehler korreliert sind, kommt es zu Verschätzungen für den Parametervektor $\mathbf{b}$ kommen. Wie zu zeigen ist, kann man das Problem korrelierender Fehler in den Griff bekommen, wenn man

$$Q(\mathbf{b}) = (\mathbf{y} - X\mathbf{b})' \Sigma^{-1} (\mathbf{y} - X\mathbf{b}) \tag{3.213}$$

minimalisiert. Man hat

$$\frac{\partial Q}{\partial \mathbf{b}} = -X' \Sigma^{-1} \mathbf{y} - (\mathbf{y} - X\mathbf{b})' \Sigma^{-1} X = \vec{0},$$

woraus

$$\hat{\mathbf{b}} = (X' \Sigma^{-1} X)^{-1} X' \Sigma^{-1} \mathbf{y} \tag{3.214}$$

folgt.

In Abschnitt 3.16.4 wird gezeigt, dass eine symmetrische, positiv-definite Matrix A in der Form $A = LL'$ dargestellt werden kann, wobei L eine untere Dreiecksmatrix ist: alle Elemente oberhalb der Diagonalen sind gleich Null; dies ist die Cholesky-Zerlegung von A . Σ ist eine symmetrische Matrix, so dass $\Sigma = LL'$ gelten muß. Dann kann (3.213) in der Form

$$Q(\mathbf{b}) = (\mathbf{y} - X\mathbf{b})' S'^{-1} S^{-1} (\mathbf{y} - X\mathbf{b}) = (L'^{-1} \mathbf{y} - S'^{-1} X\mathbf{b})' (S^{-1} \mathbf{y} - S^{-1} X\mathbf{b})$$

geschrieben werden. Dies zeigt, dass die generalisierte Kleinste-Quadrate-Schätzung äquivalent der Regression von $S^{-1}X$ auf $S^{-1}\mathbf{y}$ ist. Multipliziert man $\mathbf{y} = X\mathbf{b} + \mathbf{e}$ von links mit S^{-1} , so erhält man

$$S^{-1}\mathbf{y} = S^{-2}X\mathbf{b} + S^{-1}\mathbf{e} \quad (3.215)$$

oder

$$\tilde{\mathbf{y}} = \tilde{X}\mathbf{b} + \tilde{\mathbf{e}}, \quad \tilde{\mathbf{y}} = S^{-1}\mathbf{y}, \quad \tilde{X} = S^{-1}X. \quad (3.216)$$

Nun ist

$$\mathbb{V}(\tilde{\mathbf{e}}) = \mathbb{V}(S^{-1}\mathbf{e}) = S^{-1}\mathbb{V}(\mathbf{e})S^{-1} = S^{-1}\sigma^2 S S' S^{-1} = \sigma^2 I, \quad (3.217)$$

d.h. die Komponenten von $\tilde{\mathbf{e}}$ sind unkorreliert!

3.10.5 Extrema unter Nebenbedingungen

Es sei $f(x_1, \dots, x_n)$ eine Funktion der Variablen $x_1, \dots, x_n$. Gesucht sind die $x_{01}, \dots, x_{0n}$, für die f ein Maximum oder Minimum annimmt, wobei aber die Nebenbedingung $g(x_1, \dots, x_n) = k$, k eine Konstante, berücksichtigt werden soll, d.h. der Vektor $\mathbf{x}_0 = (x_{01}, \dots, x_{0n})'$ soll so bestimmt werden, dass auch $g(x_{01}, \dots, x_{0n}) = k$ erfüllt ist. Die Nebenbedingung kann auch in der Form $g(x_1, \dots, x_n) = 0$ angeschrieben werden.

Der Einfachheit halber wird die Extremwertbestimmung für $n = 2$ durchgeführt; das Resultat überträgt sich unmittelbar auf den Fall $n > 2$. Es wird $x = x_1, y = x_2$ gesetzt. Es soll also $f(x, y)$ unter der Nebenbedingung $g(x, y) = 0$ bestimmt werden. $g(x, y) = 0$ bedeutet, dass es eine Funktion $y = g(x)$ gibt, so dass auch $f(x, y) = f(x, g(x))$ und $g(x, g(x)) = 0$ geschrieben werden kann. Geometrisch beschreibt $f(x, y)$ eine Fläche im 3-dimensionalen Raum und $g(x, y) = 0$ beschreibt eine Kurve in der $X \times Y$ -Ebene. Die Nebenbedingung $g = 0$ bedeutet, dass man $f(x, y)$ nur für diejenigen Punkte berechnet, die auf der Kurve $g(x, y) = 0$ liegen; dafür werde $f_g = f(x, y|g(x, y) = 0)$ geschrieben. Die Menge der Punkte (x, y) , für die $f(x, y) = k$ eine Konstante ist, definiert eine Höhenlinie von $f(x, y)$. Dann existiert eine Konstante $k = c$, die die Kurve f_g genau dort berührt, wo diese ihr Maximum annimmt.

Man hat die Ableitungen

$$\frac{\partial f(x, g(x))}{\partial x} = \frac{\partial f}{\partial x} + \frac{\partial f}{\partial y} \frac{dg(x)}{dx} = f_x + f_y g',$$

wobei die Kettenregel angewendet wurde. Analog erhält man für g

$$\frac{dg}{dx} = \frac{\partial g}{\partial x} + \frac{\partial g}{\partial y} \frac{dg(x)}{dx} = g_x + g_y g'.$$

Die Extremwerte werden bestimmt, indem man die entsprechenden Ableitungen gleich Null setzt. Dementsprechend erhält man die Gleichungen

$$f_x + f_y g' = 0 \quad (3.218)$$

$$g_x + g_y g' = 0 \quad (3.219)$$

Die bisher hergeleiteten Ableitungen enthalten noch die Ableitung g' von g . Um das Extremum zu bestimmen, eliminiert man am besten g' , da die Bestimmung von g' kompliziert sein kann. Man hat $g' = -f_x/f_y = -g_x/g_y$; diese Beziehung bedeutet, dass die *Gradientenvektoren* $(f_x, f_y)'$ und $(g_x, g_y)'$ dieselbe Orientierung haben, d.h. sie unterscheiden sich allenfalls in ihrer Länge, so dass man

$$\begin{pmatrix} f_x \\ f_y \end{pmatrix} = \lambda \begin{pmatrix} g_x \\ g_y \end{pmatrix} \quad (3.220)$$

schreiben kann. $\lambda \in \mathbb{R}$ ist ein neuer, freier Parameter, der sogenannte *Lagrange-Faktor* oder auch *Lagrange-Multiplikator*. Er drückt einfach aus, dass man nur etwas über die Orientierung, nicht aber über die Länge der Gradientenvektoren am Ort des Maximums weiß. Die Vektorgleichung (3.220) zusammen mit der Bedingung $g(x, y) = 0$ führt sofort auf ein System von 3 Gleichungen in drei Unbekannten x, y und λ :

$$f_x - \lambda g_x = 0 \quad (3.221)$$

$$f_y - \lambda g_y = 0 \quad (3.222)$$

$$g(x, y) = 0 \quad (3.223)$$

Diese Überlegungen müssen nicht immer explizit durchgeführt werden, denn sie implizieren die Möglichkeit, von vorn herein die *Lagrange-Funktion* (x, y, λ) aufzustellen:

$$L(x, y, \lambda) = f(x, y) + \lambda g(x, y), \quad g(x, y) = 0 \quad (3.224)$$

L nach x, y und λ partiell zu differenzieren und die partiellen Ableitungen gleich Null zu setzen.

Damit ist das Problem der Bestimmung eines Extremums unter Nebenbedingungen gelöst. Der Parameter λ heißt *Lagrange-Faktor* oder *Lagrangescher Multiplikator* und die drei Gleichungen (3.221), (3.222) und (3.223) heißen zusammen die *Lagrangesche Multiplikatorenregel*, nach dem italo-französischen Mathematiker und Astronomen Joseph-Louis de Lagrange (1736 – 1813), der diese Regel 1788 herleitete.

Beispiel 3.14 Gegeben sei die Funktion $f(x, y) = 6 - x^2 - \frac{1}{3}y^2$ und die Nebenbedingung $x + y = 2$, die in der Form $g(x, y) = x + y - 2 = 0$ angeschrieben werden kann. Dann ist

$$f_x = -2x, \quad f_y = -\frac{2}{3}y, \quad g_x = 1, \quad g_y = 1$$

und man erhält das Gleichungssystem

$$\begin{aligned} -2x + \lambda 1 &= 0 \\ -\frac{2}{3}y + \lambda 1 &= 0 \\ x + y - 2 &= 0, \end{aligned}$$

woraus $x = 1/2$, $y = 3/2$ und $\lambda = -1$ folgt. $\square$

3.10.6 Der Rayleigh-Quotient und seine Maximierung

Bei vielen multivariaten Verfahren sollen bestimmte Größen maximiert oder minimiert werden: bei der PCA soll die Varianz der Koordinaten auf der ersten latenten Dimension maximiert werden, bei der Kanonischen Korrelation die Korrelation zwischen den latenten Variablen zweier Datensätze maximiert werden, bei der Diskriminanzanalyse soll die Varianz zwischen Gruppen von Fällen maximiert werden, etc. Es wird deshalb zunächst ein allgemeines Maximumprinzip eingeführt, das den Begriff des Rayleigh-Quotienten voraussetzt:

Definition 3.18 *Es sei A eine symmetrische Matrix. Der Quotient*

$$R(\mathbf{x}) = \frac{\mathbf{x}' A \mathbf{x}}{\mathbf{x}' \mathbf{x}} = \frac{\mathbf{x}' A \mathbf{x}}{\|\mathbf{x}\|^2} \quad (3.225)$$

heißt Rayleigh-Quotient¹⁹ oder Rayleigh-Koeffizient.

Es sei $A\mathbf{v} = \lambda\mathbf{v}$, d.h. $\mathbf{v}$ ist ein Eigenvektor von A und λ ist der zugehörige Eigenwert. Setzt man für $\mathbf{x}$ den Eigenvektor $\mathbf{v}$ in (3.225) ein, so ergibt sich

$$\frac{\mathbf{v}' A \mathbf{v}}{\mathbf{v}' \mathbf{v}} = \frac{\mathbf{v}' \lambda \mathbf{v}}{\mathbf{v}' \mathbf{v}} = \lambda \frac{\mathbf{v}' \mathbf{v}}{\mathbf{v}' \mathbf{v}} = \lambda, \quad (3.226)$$

d.h. für den Fall, dass $\mathbf{x}$ ein Eigenvektor von A ist, ist der Rayleigh-Quotient gleich dem zugehörigen Eigenwert. Der folgende Satz von Courant-Fisher verweist auf bestimmte Maximumseigenschaften von Eigenvektoren und ihren zugeordneten Eigenwerten und erweist sich als ausgesprochen nützlich bei Betrachtungen zur Bestimmung latenter Variablen wie etwa die PCA, Diskriminanzanalyse und Kanonische Korrelation.

Satz 3.21 *(Satz von Courant-Fisher) Es sei A eine symmetrische, positiv definite Matrix mit den Eigenwerten $\lambda_1 \geq \dots \geq \lambda_n$. Dann ist*

$$\max_{\mathbf{x} \neq \mathbf{0}} \frac{\mathbf{x}' A \mathbf{x}}{\mathbf{x}' \mathbf{x}} = \max_j \lambda_j = \lambda_1, \quad (3.227)$$

¹⁹Nach dem britischen Physiker John William Strutt, Dritter Baron Rayleigh (1842–1919)

und der Vektor $\mathbf{x}$, für den das Maximum angenommen wird, ist der zu λ_1 korrespondierende Eigenvektor $\mathbf{t}_1$. Weiter gilt

$$\min_{\mathbf{x} \neq \vec{0}} = \frac{\mathbf{x}' A \mathbf{x}}{\mathbf{x}' \mathbf{x}} = \min_j \lambda_j, \quad (3.228)$$

mit dem zugehörigen Eigenvektor $\mathbf{v}_{\min}$.

Zur Illustration und Übung werden zwei verschiedene Beweise gegeben.

Beweis 1: Als Nebenbedingung werde $\mathbf{x}' \mathbf{x} = \|\mathbf{x}\|^2 = 1$ gesetzt. Nach den Regeln zur Maximierung unter Nebenbedingungen ist dann die Funktion

$$Q = \frac{\mathbf{x}' A \mathbf{x}}{\mathbf{x}' \mathbf{x}} - \lambda(\mathbf{x}' \mathbf{x} - 1) = \mathbf{x}' A \mathbf{x} - \lambda(\mathbf{x}' \mathbf{x} - 1)$$

zu maximieren, wobei λ ein Lagrange-Faktor ist. Man erhält sofort

$$\frac{\partial Q}{\partial \mathbf{x}} = 2A\mathbf{x} - 2\lambda\mathbf{x}.$$

Es sei $\partial Q / \partial \mathbf{x} = 0$ für $\mathbf{x} = \mathbf{u}$. Dann folgt sofort $A\mathbf{u} = \lambda\mathbf{u}$, d.h. der Rayleigh-Quotient wird maximal, wenn $\mathbf{x} = \mathbf{u}$ gilt, wenn also $\mathbf{x}$ gleich dem ersten Eigenvektor von A ist, und wenn λ gleich dem zugehörigen Eigenwert ist. $\square$

Beweis 2: Es sei $A = T\Lambda T'$ die Spektralzerlegung von A , d.h. T sei die Matrix der Eigenvektoren von A und Λ die dazu korrespondierende Diagonalmatrix der Eigenwerte von A . Dann hat man

$$\frac{\mathbf{x}' A \mathbf{x}}{\mathbf{x}' \mathbf{x}} = \frac{\mathbf{x}' T \Lambda T' \mathbf{x}}{\mathbf{x}' T T' \mathbf{x}},$$

denn $TT' = I$. Mit $\mathbf{s} = T' \mathbf{x}$ erhält man dann

$$\frac{\mathbf{x}' T \Lambda T' \mathbf{x}}{\mathbf{x}' T T' \mathbf{x}} = \frac{\mathbf{s}' \Lambda \mathbf{s}}{\mathbf{s}' \mathbf{s}} = \frac{\sum_{j=1}^n \lambda_j s_j^2}{\sum_{j=1}^n s_j^2} \leq \lambda_{\max},$$

denn wenn man die λ_j durch $\lambda_{\max} = \max_j \lambda_j$ ersetzt, ist ja

$$\sum_{j=1}^n \lambda_j s_j^2 \leq \sum_{j=1}^n \lambda_{\max} s_j^2 = \lambda_{\max} \sum_{j=1}^n s_j^2, \quad (3.229)$$

und die Summe der s_j^2 kürzt sich heraus. Ersetzt man die λ_j durch $\min_j \lambda_j$, so findet man

$$\sum_{j=1}^n \lambda_j s_j^2 \geq \min_j \lambda_j \sum_{j=1}^n s_j^2, \quad (3.230)$$

und man hat (3.228) nachgewiesen.

Jetzt muß noch gezeigt werden, für welchen Vektor $\mathbf{x}$ das Maximum angenommen wird. Es sei $\mathbf{t}_1$ der zu $\lambda_{\max} = \lambda_1$ korrespondierende Eigenvektor. Dann werde $\mathbf{x} = \mathbf{t}_1$ gesetzt. Dann folgt $\mathbf{s} = T'\mathbf{x} = T'\mathbf{t}_1$, und wegen der Orthonormalität der Eigenvektoren in T erhält man $T'\mathbf{t}_1 = (1, 0, \dots, 0)' = \mathbf{e}_1$, so dass $\mathbf{s}'\mathbf{s} = \mathbf{e}_1'\mathbf{e}_1 = 1$, und

$$\|A\| = \max_{\mathbf{x} \neq 0} \frac{\|A\mathbf{x}\|_2}{\|\mathbf{x}\|_2} = \frac{\|A\mathbf{t}_1\|_2}{\mathbf{t}_1'\mathbf{t}_1} = \mathbf{t}_1' A \mathbf{t}_1 = \lambda_1.$$

Für $\mathbf{x} = \mathbf{t}_1$ wird also der Maximalwert λ_1 angenommen. $\square$

Beispiel 3.15 Latente Variable mit maximaler Varianz: Es sei X eine $(m \times n)$ -Datenmatrix; die SVD für X ist $X = Q\Lambda^{1/2}P' = LP'$. L enthält orthogonale Basisvektoren für die Spaltenvektoren von X . Es ist $XX' = Q\Lambda Q'$, Q die Matrix der orthonormalen Eigenvektoren von XX' , so dass $Q'(XX')Q = \Lambda$, d.h. $\mathbf{q}_1'(XX')\mathbf{q}_1 = \lambda_1$, und bei der üblichen Aordnung $\lambda_1 \geq \dots \geq \lambda_n$ ist dies der maximale Wert des Rayleigh-Quotienten. Nach (3.188) ist für zentrierte Messwerte die Varianz der Komponenten des k -ten Spaltenvektors $\mathbf{L}_k$ von L durch $\mathbf{L}_k'\mathbf{L}_k = \|\mathbf{L}_k\|^2 = \lambda_k$ gegeben (bis auf den Faktor $1/m$). Nach Satz 3.21 haben dann die Komponenten von $\mathbf{L}_1$ die größtmögliche Varianz, die von $\mathbf{L}_2$ haben die zweitmaximale Varianz, etc. Denn mit $A = XX'$ gilt ja $A = Q\Lambda Q'$ und

$$\max_{\mathbf{x}} \frac{\mathbf{x}'A\mathbf{x}}{\mathbf{x}'\mathbf{x}} = \mathbf{q}_1'A\mathbf{q}_1 = \lambda_1,$$

und $\mathbf{L}_1 = \lambda_1\mathbf{q}_1$, $\mathbf{L}_1'\mathbf{L}_1 = \lambda_1\mathbf{q}_1'\mathbf{q}_1 = \lambda_1$ wegen der Normiertheit der $\mathbf{q}_k$. Dies ist das PCA-Prinzip. $\square$

Die folgende Aussage vervollständigt den Satz von Courant-Fischer und ist eher ein Korollar zu diesem Satz.

Satz 3.22 *Es sei A wie in Satz 3.21 definiert. Dann gilt*

$$\max_{\mathbf{x} \perp \mathbf{t}_1, \dots, \mathbf{t}_k} \frac{\mathbf{x}'A\mathbf{x}}{\mathbf{x}'\mathbf{x}} = \lambda_{k+1}, \quad k < n, \quad (3.231)$$

für $\mathbf{x} = \mathbf{t}_{k+1}$ der $(k+1)$ -te Eigenvektor. ($\perp$ steht für "ist orthogonal zu".)

Beweis: Es sei wieder $T = [\mathbf{t}_1 | \dots | \mathbf{t}_n]$ die Matrix der Eigenvektoren von A . Dann existieren reelle Zahlen $y_1, \dots, y_n$ derart, dass ein Vektor $\mathbf{x}$ in der Form $\mathbf{x} = T\mathbf{y}$ dargestellt werden kann, wobei die Komponenten von $\mathbf{y}$ durch die y_j gegeben sind. Nun soll speziell $\mathbf{x} \perp \mathbf{t}_1, \dots, \mathbf{t}_k$ gelten. Dann muß aber

$$\mathbf{v}_k'\mathbf{x} = y_1\mathbf{t}_k'\mathbf{t}_1 + y_2\mathbf{t}_k'\mathbf{t}_2 + \dots + y_n\mathbf{t}_k'\mathbf{t}_n = y_k = 0$$

gelten, denn $\mathbf{t}_k'\mathbf{t}_k = 1$, so dass $y_k\mathbf{t}_k'\mathbf{t}_k = y_k$. Die Forderung der Orthogonalität von $\mathbf{x}$ zu den ersten k Eigenvektoren impliziert also $y_1 = \dots = y_k = 0$. Dann folgt aus (3.227)

$$\frac{\mathbf{x}'A\mathbf{x}}{\mathbf{x}'\mathbf{x}} = \frac{\sum_{j=k+1}^n \lambda_j y_j^2}{\sum_{j=k+1}^n y_j^2},$$

und analog zur Argumentation im Beweis zu Satz 3.21 folgt (3.231). $\square$

Anmerkung: Der Satz 3.22 ist hier als eine Art Korollar zu Satz 3.21 aufgeführt worden, aber Satz 3.22 kann natürlich auch als die allgemeine Version des Satzes von Courant-Fischer betrachtet werden: Für $k = 0$ betrachtet man $\max_{\mathbf{x} \neq \mathbf{0}} \frac{\mathbf{x}' A \mathbf{x}}{\mathbf{x}' \mathbf{x}} = \lambda_1$, und für $k = n - 1$ erhält man

$$\max_{\mathbf{x} \perp \mathbf{t}_1, \dots, \mathbf{t}_{n-1}} \frac{\mathbf{x}' A \mathbf{x}}{\mathbf{x}' \mathbf{x}} = \lambda_n.$$

$\square$

3.10.7 Vektor- und Matrixnormen

Es ist immer wieder von normierten Vektoren die Rede gewesen: ein Vektor $\mathbf{x}$ ist normiert, wenn $\|\mathbf{x}\| = 1$, wobei $\|\mathbf{x}\|$ die Länge im Sinne des Satzes von Pythagoras ist, man spricht auch von Euklidischer Norm. Dies ist ein Spezialfall, die Norm eines Vektors kann allgemeiner definiert werden.

Die Norm eines Vektors definiert, in welchem Sinne von der "Größe" eines Vektors gesprochen werden soll, – die übliche euklidische Norm $\|\mathbf{x}\| = (\sum_i x_i^2)^{1/2}$ definiert die Länge des Vektors $\mathbf{x}$ als seine "Größe"²⁰. Ebenso kann eine *Matrixnorm* definiert werden. Dieser Begriff erweist sich als nützlich, wenn bestimmte Maxima oder Minima gefunden werden sollen, etwa die Varianzen von Projektionen einer Punktekongfiguration auf bestimmte Dimensionen, oder die Güte der Approximation an eine Datenmatrix. Es wird zuerst der Begriff der Vektornorm spezifiziert:

Definition 3.19 *Es seien $\mathbf{x} \in \mathbb{R}^n$ n -dimensionale Vektoren. Eine Vektornorm ist eine Abbildung $f : \mathbb{R}^n \rightarrow \mathbb{R}$ (d.h. es wird einem Vektor $\mathbf{x}$ eine bestimmte reelle Zahl zugeordnet), die den Bedingungen*

1. $f(\mathbf{x}) \geq 0$,
2. $f(\mathbf{x} + \mathbf{y}) \leq f(\mathbf{x}) + f(\mathbf{y})$,
3. $f(a\mathbf{x}) = a f(\mathbf{x})$, für $a \in \mathbb{R}$

genügt. Dann heißt f eine Vektornorm. f wird durch $f(\mathbf{x}) = \|\mathbf{x}\|$ notiert. Der Einheitsvektor in Bezug auf eine Norm $\|\cdot\|$ ist derjenige Vektor, für den $\|\mathbf{x}\| = 1$ gilt.

Von besonderem Interesse sind die p -Normen

$$\|\mathbf{x}\|_p = (|x_1|^p + \dots + |x_n|^p)^{1/p} = \left(\sum_{j=1}^n |x_j|^p \right)^{1/p}. \quad (3.232)$$

²⁰Um sich eine inhaltliche Vorstellung zu machen, stelle man sich vor dass die Komponenten x_i von $\mathbf{x}$ Maße für Begabungen $M_1, \dots, M_n$ repräsentieren. Dann ist $\|\mathbf{x}\|$ ein *mögliches* Maß für die Gesamtbegabung einer Person.

Für $p = 1$ erhält man die 1-Norm

$$\|\mathbf{x}\|_1 = (|x_1| + \cdots + |x_n|) = \sum_{j=1}^n |x_j|. \quad (3.233)$$

und für $p = 2$ die euklidische Norm

$$\|\mathbf{x}\|_2 = (|x_1|^2 + \cdots + |x_n|^2)^{1/2} = \left(\sum_{j=1}^n |x_j|^2 \right)^{1/2} = \sqrt{\mathbf{x}'\mathbf{x}}. \quad (3.234)$$

Für $p = \infty$ schließlich findet man die Maximum-Norm

$$\|\mathbf{x}\|_\infty = \max_{1 \leq i \leq n} |x_i|. \quad (3.235)$$

Die Maximum-Norm ergibt sich aus der p -Norm für $p \rightarrow \infty$. Es sei $x_k = x_{\max}$ die maximale Komponente von $\mathbf{x}$. Dann ist

$$\|\mathbf{x}\|_p = \left(|x_k|^p \sum_{j=1}^n \frac{|x_j|^p}{|x_k|^p} \right)^{1/p}.$$

Wegen $|x_j|/|x_k| \leq 1$ für alle $j \neq k$ folgt $\lim_{p \rightarrow \infty} |x_j|^p/|x_k|^p \rightarrow 0$ für $j \neq k$ und $|x_j|^p/|x_k|^p = 1$ für $j = k$, so dass

$$\lim_{p \rightarrow \infty} \|\mathbf{x}\|_p = x_k = x_{\max}.$$

Matrixnormen: Der Begriff der Norm kann auch auf Matrizen angewendet werden:

Definition 3.20 *Es sei $\mathbb{R}^{m \times n}$ die Menge der reellen $(m \times n)$ -Matrizen²¹. Eine Matrixnorm ist eine Abbildung $\|\cdot\| : \mathbb{R}^{m \times n} \rightarrow \mathbb{R}_+$, $\mathbb{R}_+$ die Menge der reellen Zahlen größer oder gleich Null, und $A \mapsto \|A\|$ derart, dass*

1. $\|A\| = 0$ genau dann, wenn $A = 0$ die Nullmatrix ist,
2. $\|\lambda X\| = \lambda \|A\|$,
3. $\|A + B\| \leq \|A\| + \|B\|$

gilt. Zusammen mit der Norm $\|\cdot\|$ wird der Vektorraum der $(m \times n)$ -Matrizen dann zu einem normierten Vektorraum $(\mathbb{R}^{m \times n}, \|\cdot\|)$.

Es gibt verschiedene Normen, von denen hier einige als Beispiel genannt werden:

²¹Diese Definition ist etwas vereinfacht formuliert, eigentlich muß es heißen: es sei $\mathbb{K} = \mathbb{R}$ der Körper der reellen Zahlen und $\mathbb{K}^{m \times n} = \mathbb{R}^{m \times n}$ die Menge der reellen $(m \times n)$ -Matrizen, etc

1. Die *Frobenius-Norm*.

$$\|A\|_F = \left(\sum_{i=1}^m \sum_{j=1}^n |a_{ij}|^2 \right)^{1/2}. \quad (3.236)$$

Für diese Norm wird auch der Name *Schur-Norm* oder *Hilbert-Schmidt-Norm* verwendet.

2. Die p -Norm: sie ist definiert durch

$$\|A\|_p = \max_{\mathbf{x} \neq \vec{0}} \frac{\|A\mathbf{x}\|_p}{\|\mathbf{x}\|_p} \quad (3.237)$$

Da $A\mathbf{x} = \mathbf{y}$ ein Vektor ist, ist $\|A\mathbf{x}\|_p$ eigentlich eine Vektornorm; allgemein heißen Normen der Form

$$\|A\| = \max_{\mathbf{x} \neq \vec{0}} \frac{\|A\mathbf{x}\|}{\|\mathbf{x}\|}. \quad (3.238)$$

durch eine Vektornorm induzierte Normen.

Die Frobenius- und die p -Norm sind die am häufigsten vorkommenden Matrixnormen. Für $\|A\|_p$ gilt, wenn A eine $(m \times n)$ -Matrix ist,

$$\|A\|_p = \sup_{\mathbf{x} \neq \vec{0}} \left\| \left(A \frac{b\mathbf{x}}{\|\mathbf{x}\|} \right) \right\|_p = \max_{\|\mathbf{x}\|_p=1} \|A\mathbf{x}\|_p. \quad (3.239)$$

Speziell für $p = 2$ ist mit $\mathbf{y} = A\mathbf{x}$ die Norm $\|A\mathbf{x}\|_2$ durch die Norm $\|\mathbf{y}\|_2 = (\mathbf{y}'\mathbf{y})^{1/2} = \|\mathbf{y}\|$ gegeben, und nach dem Courant-Fischer Theorem 3.21 findet man

$$\|A\|_2 = \max_{\|\mathbf{x}\|_2=1} \|A\mathbf{x}\|_2 = \sqrt{\lambda_{\max}}, \quad (3.240)$$

wobei $\lambda_{\max}$ der maximale Eigenwert von $A'A$ ist. Für die Frobenius-Norm findet man

Satz 3.23 *Es sei A eine $(m \times n)$ -Matrix. Für die Frobenius-Norm $\|A\|_F$ gilt*

$$\|A\|_F^2 = \text{spur}(AA') = \sum_{i=1}^n \lambda_i \quad (3.241)$$

wobei $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$ die Eigenwerte von $A'A$ sind.

Beweis: Auf A kann die SVD angewendet werden: $A = Q\Lambda^{1/2}P'$, $\lambda_j \geq 0$ für $j = 1, \dots, n$. Dann ist $AA' = Q\Lambda^{1/2}P'P\Lambda^{1/2}Q' = Q\Lambda Q'$, und die Diagonalelemente

von $Q\Lambda Q'$ sind von der Form $\sum_i \lambda_i q_{ji}^2$ für $j = 1, \dots, n$. Die Spur von $A'A$ ist die Summe dieser Diagonalelemente, d.h.

$$\text{spur}(AA') = \sum_{j=1}^n \sum_{i=1}^n \lambda_i q_{ji}^2 = \sum_{i=1}^n \lambda_i \sum_{j=1}^n q_{ij}^2.$$

Aber $\sum_{j=1}^n q_{ij}^2 = 1$ für alle i , da Q orthonormal ist, d.h. die Eigenvektoren haben alle die Länge 1. Damit ist (3.241) gezeigt. $\square$

Anmerkung: $A = Q\Lambda^{1/2}P'$ impliziert $A'A = P\Lambda P'$ und wegen der Orthonormalität der Spaltenvektoren von P folgt in analoger Weise

$$\text{spur}(A'A) = \sum_{i=1}^n \lambda_i. \quad (3.242)$$

3.10.8 Die Approximation von Matrizen

Der folgende Satz macht eine Aussage über die Güte der Approximation einer Matrix A durch eine Matrix mit kleinerem Rang. So sei etwa $A = X$ eine Datenmatrix mit dem Rang n und man will versuchen, X durch eine Matrix X_r mit dem Rang $r < n$ zu approximieren, d.h. durch möglichst wenige latente Variable zu "erklären".

Satz 3.24 *Es seien A und A_k ($m \times n$)-Matrizen, $m \geq n$, und es seien die Matrizen A und A_k durch*

$$A = Q\Lambda^{1/2}P' = \sum_{j=1}^n \sqrt{\lambda_j} \mathbf{q}_j \mathbf{p}_j', \quad A_k = Q\Lambda_k^{1/2}P' = \sum_{j=1}^k \sqrt{\lambda_j} \mathbf{q}_j \mathbf{p}_j' \quad (3.243)$$

definiert, wobei $\Lambda = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_n)$ mit $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$ und $\Lambda_k = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_k)$ mit $k < n$ sei. Dann gilt

$$\|A - A_k\|_2 = \sqrt{\lambda_{k+1}}. \quad (3.244)$$

Beweis: Es ist $A = Q\Sigma P'$, $A_k = Q\Sigma_k P'$, wobei $\Sigma = \Lambda^{1/2}$, $\Sigma_k = \Lambda_k^{1/2}$, Λ die Diagonalmatrix der Eigenwerte von $A'A$, Λ_k die Diagonalmatrix der ersten k Eigenwerte. Dann ist

$$A - A_k = Q\Sigma P' - Q\Sigma_k P' = Q(\Sigma - \Sigma_k)P' = Q\Sigma^* P',$$

$\Sigma^* = \text{diag}(\underbrace{0, \dots, 0}_k, \sigma_{k+1}, \dots, \sigma_n)$, $\sigma_j = \sqrt{\lambda_j}$. Da $\|A\|_2 = \sigma_{\max} = \sigma_1$ im Falle geordneter Singularwerte $\sigma_{k+1} \geq \dots \geq \sigma_n$, folgt

$$\|A - A_k\|_2 = \sigma_{k+1} = \sqrt{\lambda_{k+1}},$$

da nun σ_{k+1} der maximale Singularwert ist. $\square$

Anmerkungen:

1. Die Approximation wird trivialerweise immer besser, je größer der Wert von k , da ja der Wert von λ_{k+1} mit größer werdendem k immer kleiner wird. Der nichttriviale Teil der Aussage ist, dass $\|A - A_k\|_2$ gerade dem Wert von $\sqrt{\lambda_{k+1}}$ entspricht.
2. Bei der Approximation von A durch A_k wurde von der SVD von A Gebrauch gemacht. Die Gleichung $A = Q\Lambda^{1/2}P'$ ist insofern trivial, als die SVD stets gilt. Dass man A durch A_k approximiert, wobei A_k nur durch die ersten k Terme der SVD definiert ist, kann zur Frage führen, ob es eine andere Repräsentation für A_k gibt, die nicht auf der SVD beruht, aber besser ist in dem Sinne, dass $\|A - A_k\|_2 < \sqrt{\lambda_{k+1}}$. Eine solche gibt es nicht, wie noch gezeigt werden wird.
3. Man vergleiche die Aussage (3.244) mit der Aussage (3.243) von Satz 3.21. Wie die Gleichung (3.181), also die SVD von A , zeigt, ist A additiv durch Matrizen aufgebaut, die jeweils als dyadisches Produkt der Singulärvektoren $\mathbf{q}_j$ und $\mathbf{p}_j$ definiert sind und die jeweils den Rang 1 haben (vergl. Satz 3.4, Seite 82). Der Rang $\text{rg}(A) \leq \min(m, n)$ ist durch die Anzahl der von Null verschiedenen Eigenwerte λ_j und damit durch die Anzahl der von Null verschiedenen $\sigma_k \mathbf{q}_k \mathbf{p}_k$ gegeben. Da die ersten k Eigenwerte von A und A_k identisch sind, enthält die Differenz $\Lambda^{1/2} - \Lambda_k^{1/2}$ nur Nullen, und der erste von Null verschiedene Wert in der Diagonalen ist $\sigma_{k+1} = \sqrt{\lambda_{k+1}}$. Da die Eigenwerte λ_j in Λ der Größe nach angeordnet sind, ist σ_{k+1} nun der größte Singulärwert für $A - A_k$. $\square$

Im Folgenden bedeutet $\min_{\text{rg}(B)=k} \|X - B\|$ bzw. $\min_{\text{rg}(B)=k} \|X - B\|_F$ diejenige Matrix B , die (i) den Rang $\text{rg}(B) = k$ hat und die (ii) den Wert für die Norm $\|X - B\|$ bzw. $\|X - B\|_F$ minimiert. Es kann nun der folgende Satz bewiesen werden:

Satz 3.25 *Es seien A und B $(m \times n)$ -Matrizen mit $m > n$, wobei A den Rang r und B den Rang $k < r$ habe. Weiter sei*

$$A_k = Q\Lambda_k^{1/2}P' = \sum_{j=1}^k \sqrt{\lambda_j} \mathbf{q}_j \mathbf{p}'_j, \quad (3.245)$$

Λ_k die zu den ersten²² k Eigenvektoren korrespondierenden Eigenwerte von X enthält. Dann gilt

$$\min_{B \in \mathbb{R}^{m,n}, \text{rg}(B)=k} \|A - B\|_2 = \|A - A_k\|_2 = \sigma_{k+1} = \sqrt{\lambda_{k+1}} \quad (3.246)$$

²²Es wird angenommen, dass die Eigenwerte der Größe nach geordnet sind, $\lambda_1 > \lambda_2 > \dots > \lambda_k$

Anmerkung: Dieser Satz wird gelegentlich als Satz von Eckart & Young bezeichnet, weil Eckart & Young (1936) eine derartige Aussage vorgestellt haben, allerdings nicht mit diesem Beweis. Tatsächlich hat schon Schmidt (1907) diese Aussage vorgestellt, und Mirsky (1960) hat diesen und den folgenden Satz 3.26 in allgemeiner Weise bewiesen, so dass auch zusammenfassend vom Schmidt-Mirsky-Theorem gesprochen wird.

Beweis: Zur Vereinfachung werde

$$\|A - B_{\min}\| = \min_{B \in \mathbb{R}^{m,n}, \text{rg}(B)=k} \|A - B\|_2$$

gesetzt. Zu zeigen ist, dass

$$\min_{B \in \mathbb{R}^{m,n}, \text{rg}(B)=k} \|A - B\|_2 = \|A - A_k\|_2.$$

Dazu werde angenommen, dass

$$\|A - B_{\min}\| < \|A - A_k\|_2.$$

Die Ungleichung bleibt bestehen, wenn beide Seiten mit dem gleichen Faktor (> 0) multipliziert werden. Für alle n -dimensionalen Vektoren $\mathbf{b}$ gilt dann

$$\|A - B_{\min}\| \|\mathbf{b}\| < \|A - A_k\|_2 \|\mathbf{b}\| = \sigma_{k+1} \|\mathbf{b}\|.$$

Dann gilt auch

$$\|(A - B_{\min})\mathbf{b}\| \leq \|(A - A_k)\mathbf{b}\| \leq \sigma_{k+1} \|\mathbf{b}\|.$$

Insbesondere kann dann $\mathbf{b}$ als Linearkombination der ersten $k+1$ (Eigen-)Vektoren von P gewählt werden: sind also gerade die ersten $k+1$ Spalten von P die Spalten von P_{k+1} , so sei $\mathbf{b} = P_{k+1}\mathbf{x} = x_1\mathbf{p}_1 + \cdots + x_{k+1}\mathbf{p}_{k+1}$. Es ist

$$B_{\min}\mathbf{b} = B_{\min}P_{k+1}\mathbf{x},$$

$B_{\min}P_{k+1}$ ist also eine $(m \times (k+1))$ -Matrix. Nach Satz 3.3, Gleichung (3.69) (Seite 82) ist aber $\text{rg}(B_{\min}P_{k+1}) \leq \min(\text{rg}(B_{\min}), \text{rg}(P_{k+1})) = k$, da ja $\text{rg}(B_{\min}) = k$ nach Voraussetzung. Dann folgt aber

$$\text{rg}(B_{\min}P_{k+1}) + \dim(\text{kern}(B_{\min}P_{k+1})) = k + 1,$$

d.h.

$$\dim(\text{kern}(B_{\min}P_{k+1})) \geq k + 1 - k = 1.$$

Also enthält $\text{kern}(B_{\min}P_{k+1})$ mindestens einen Vektor $\mathbf{x}$ mit $B_{\min}P_{k+1}\mathbf{x} = \vec{0}$. Es sei also $\mathbf{x} \in \text{kern}(B_{\min}P_{k+1})$; dann folgt

$$\|A\mathbf{b} - B_{\min}\mathbf{b}\| = \|AP_{k+1}\mathbf{x}\| < \|(AP_{k+1}\mathbf{x} - A_kP_{k+1}\mathbf{x})\| \leq \sigma_{k+1} \|\mathbf{b}\|.$$

Es ist aber

$$\|A\| \|P_{k+1} \mathbf{x}\| < \|(A - A_k P_{k+1}) \mathbf{x}\| \leq \sigma_{k+1} \|\mathbf{b}\|,$$

d.h.

$$\sigma_1 < \sigma_{k+1},$$

im Widerspruch zu $\sigma_1 \geq \sigma_{k+1}$. Damit gilt (3.246). $\square$

Satz 3.26 (Satz von Schmidt-Mirsky) *Es sei A und $B_{\min}$ ($m \times n$)-Matrizen mit $m > n$, wobei A den Rang r und B den Rang $k < r$ habe, $B_{\min}$ sei wie in Satz 3.25 definiert und A_k sei wie in (3.245) definiert. Dann gilt*

$$\|A - B_{\min}\|_F = \sqrt{\sum_{j=k+1}^n \lambda_j}, \quad (3.247)$$

wobei $\|\cdot\|_F$ die Frobenius-Norm ist.

Beweis: Die Anwendung der SVD auf $A - A_k$ liefert

$$\|A - A_k\|^2 = \|Q(\Lambda^{1/2} - \Lambda_k^{1/2})P'\|^2 = \sum_{j=1}^n \lambda_j - \sum_{j=k+1}^n \lambda_j = \|A\|_F^2 - \sum_{j=k+1}^n \lambda_j.$$

$B_{\min}$ kann als Summe von durch dyadische Produkte definierte Matrizen definiert werden, also

$$B_{\min} = \sum_{j=1}^n \mathbf{x}_j \mathbf{y}_j',$$

wobei die $\mathbf{x}_j$ m -dimensionale und die $\mathbf{y}_j$ n -dimensionale Vektoren sind. Da auch für $B_{\min}$ eine Singularwertzerlegung gilt, können für $\mathbf{x}_j$ und $\mathbf{y}_j$ die jeweils mit $\sqrt{\sigma_j}$ multiplizierten Links- und Rechtssingulärvektoren von $B_{\min}$ gewählt werden, d.h. man kann orthogonale Vektoren wählen. Zu zeigen ist dann, dass

$$\|A - \sum_{j=1}^k \mathbf{x}_j \mathbf{y}_j\| \geq \|A\|^2 - \sum_{j=1}^k \lambda_j.$$

Nach Definition der Frobenius-Norm hat man

$$\begin{aligned} \|A - \sum_{j=1}^k \mathbf{x}_j \mathbf{y}_j\|_F^2 &= \text{spur} \left((A - \sum_{j=1}^k \mathbf{x}_j \mathbf{y}_j)' (A - \sum_{j=1}^k \mathbf{x}_j \mathbf{y}_j) \right) \\ &= \text{spur} \left(A' A + \sum_{j=1}^k (\mathbf{y}_j - A' \mathbf{x}_j)(\mathbf{y}_j - A' \mathbf{x}_j)' - \sum_{j=1}^k A' \mathbf{x}_j \mathbf{x}_j A \right) \end{aligned}$$

Es ist $\text{spur}((\mathbf{y}_j - A'\mathbf{x}_j)(\mathbf{y}_j - A'\mathbf{x}_j)) \geq 0$, $\text{spur}(A'\mathbf{x}_j\mathbf{x}_j'A) = \|A'\mathbf{x}_j\|^2$ und es ist zu zeigen, dass

$$\sum_{j=1}^k \|A'\mathbf{x}_j\|^2 \leq \sum_{j=1}^k \lambda_j.$$

Die SVD von A sei $A = Q\Sigma P'$, und es sei $P_1 = [|\mathbf{p}_1\rangle \dots |\mathbf{p}_k\rangle]$, $P_2 = [|\mathbf{p}_{k+1}\rangle \dots |\mathbf{p}_n\rangle]$, so dass $P = [P_1|P_2]$. Analog dazu sei $\Sigma_1 = \text{diag}(\sigma_1, \dots, \sigma_k)$, $\Sigma_2 = \text{diag}(\sigma_{k+1}, \dots, \sigma_n)$. Dann hat man

$$\begin{aligned} \|A'\mathbf{x}_j\|^2 &= \|Q\Sigma P'\mathbf{x}_j\|^2 = \|\Sigma P'\mathbf{x}_j\|^2 = \\ &= \|\Sigma_1 P'_1 \mathbf{x}_j\|^2 + \|\Sigma_2 P'_2 \mathbf{x}_j\|^2 = \lambda_k + \lambda_k + \lambda_k (\|P'\mathbf{x}_j\|^2 - \|P'_1 \mathbf{x}_j\|^2 - \|P'_2 \mathbf{x}_j\|^2) \\ &= \lambda_k + (\underbrace{\|\Sigma_1 P'_1 \mathbf{x}_j\|^2}_{(1)} - \underbrace{\lambda_k \|P'_1 \mathbf{x}_j\|^2}_{(2)}) \end{aligned}$$

Der Term (1) ist positiv, ebenso (2), da P orthonormal, und $\mathbf{x}_j$ ist ebenfalls orthonormal. Dann folgt

$$\begin{aligned} \sum_{j=1}^k \|A'\mathbf{x}_j\|^2 &\leq k\lambda_k + \sum_{j=1}^k (\|\Sigma_1 P'_1\|^2 - \lambda_k \|P'_1 \mathbf{x}_j\|^2) \\ &= k\lambda_k + \sum_{j=1}^k \sum_{i=1}^k (\lambda_i - \lambda_j) |\mathbf{v}_j' \mathbf{x}_j|^2 \\ &\leq \sum_{i=1}^k (\lambda_k + (\lambda_i - \lambda_k)) = \sum_{j=1}^k \lambda_j. \end{aligned}$$

□

3.11 Basen und Transformationen von Basen

Es sei V ein n -dimensionaler Vektorraum. Jede Menge von n linear unabhängigen, n -dimensionalen Vektoren aus V ist eine Basis von V , und es fragt sich, welche Basis man wählen soll, wenn man eine Menge von Datenvektoren $\mathbf{x}_j$, $j = 1, \dots, n$ als Linearkombination von Basisvektoren repräsentieren will. Bevor auf diese Frage näher eingegangen wird, sollen ein paar grundsätzliche Sachverhalte geklärt werden.

Es seien $\mathcal{B} = (\mathbf{b}_1, \dots, \mathbf{b}_n)$ und $\mathcal{C} = (\mathbf{c}_1, \dots, \mathbf{c}_n)$ irgend zwei Basen des Vektorraums V . Dann gilt

$$\mathcal{L}(\mathcal{B}) = \mathcal{L}(\mathcal{C}) = V$$

und damit sind die $\mathbf{b}_j \in \mathcal{L}(\mathcal{C})$ und die $\mathbf{c}_j \in \mathcal{L}(\mathcal{B})$. Jeder Basisvektor $\mathbf{c}_j$ lässt sich als Linearkombination der Basisvektoren $\mathbf{b}_j$ darstellen und umgekehrt, d.h. es gelten die Gleichungen

$$\mathbf{c}_j = t_{1j}\mathbf{b}_1 + \dots + t_{nj}\mathbf{b}_n = B\mathbf{t}_j, \quad j = 1, \dots, n \quad (3.248)$$

$$\mathbf{b}_j = s_{1j}\mathbf{c}_1 + \dots + s_{nj}\mathbf{c}_n = C\mathbf{s}_j \quad (3.249)$$

wobei B eine Matrix ist, die entsteht, wenn man die $\mathbf{b}_j$ spaltenweise zu einer Matrix zusammenfasst, und C ist analog die Matrix, deren Spaltenvektoren die $\mathbf{c}_j$ sind. $\mathbf{t}_j = (t_{1j}, \dots, t_{nj})'$ und $\mathbf{s}_j = (s_{1j}, \dots, s_{nj})'$. Ist T die Matrix mit den $\mathbf{t}_j$ als Spaltenvektoren und S die Matrix mit den $\mathbf{s}_j$ als Spaltenvektoren, so sind (3.248) und (3.249) äquivalent zu

$$C = BT \quad (3.250)$$

$$B = CS. \quad (3.251)$$

Setzt man die rechte Seite von (3.251) für B in (3.250) ein, so erhält man $C = CST$. C hat vollen Rang, so dass C^{-1} existiert und es folgt $C^{-1}C = C^{-1}CST = ST = I$. Da C und B Basen repräsentieren, müssen beide Matrizen den vollen Rang n haben, so dass $\text{rg}(C) = \text{rg}(BT)$. Nach Satz 3.3, Seite 82, gilt

$$\text{rg}(C) \leq \min(\text{rg}(B), \text{rg}(T)) \leq \text{rg}(B).$$

Aus der Definition von B und C als Matrizen, deren Spalten Basisvektoren sind, folgt natürlich schon, dass $\text{rg}(B) = \text{rg}(C)$ und damit $\text{rg}(B) = \text{rg}(C) = \text{rg}(T)$, woraus wiederum die Existenz von T^{-1} folgt. Dann hat man aber $STT^{-1} = T^{-1}$, also

$$S = T^{-1}. \quad (3.252)$$

Analog folgt die Existenz von S^{-1} .

Formal sind die verschiedenen Basen eines Vektorraums einander äquivalent. Will man Daten interpretieren, so muß man sich für eine der Basen entscheiden. Die Entscheidung muß anhand von Kriterien geschehen, die sich aus dem Begriff des Vektorraums oder dem der Basis eines Vektorraumes selbst nicht herleiten lassen. Hier soll noch auf die Suche nach einer Basis für eine gegebene $(m \times n)$ -Datenmatrix X eingegangen werden.

Formal gesehen sind die Spaltenvektoren von X eine Stichprobe von n m -dimensionalen Vektoren aus einem m -dimensionalen Vektorraum, bzw. eine Stichprobe von m n -dimensionalen Vektoren aus einem n -dimensionalen Vektorraum. Für $m > n$ kann nur eine Teilbasis mit maximal n m -dimensionalen Basisvektoren gewählt werden. Solche Teilbasen sollen mit $\mathcal{B}_m^n$ bezeichnet werden: der untere Index m gibt die Dimensionalität der Vektoren an, der obere Index n mit $n < m$ die Anzahl der linear unabhängigen Vektoren, die in der Teilbasis enthalten sein müssen. Natürlich gibt es viele solche Teilbasen, aber man muß aus der Menge dieser Teilbasen eine Basis $\mathcal{B}_m^n$ auswählen, deren lineare Hülle die Datenvektoren enthält: $\mathbf{x}_1, \dots, \mathbf{x}_n \in \mathcal{L}(\mathcal{B}_m^n)$. Dass nicht jede Teilbasis $\mathcal{B}_m^n$ die Datenvektoren erklären kann macht man sich anschaulich klar, wenn man an die Menge der 2-dimensionalen Teilräume – also die Menge der Ebenen – in einem 3-dimensionalen Raum denkt: wenn zwei Ebenen nicht parallel sind, schneiden sich und die Schnittfläche ist eine Gerade, also ein 1-dimensionaler Teilraum. Zwei Basen, von denen die eine die Vektoren aus der Ebene E_1 generiert und die zweite die aus der

Ebene E_2 können 3-dimensionale Datenvektoren nur dann als Linearkombinationen erzeugen, wenn die Datenvektoren eben in diesem 1-dimensionalen Teilraum liegen.

Man wird also nur solche Basen wählen können, die in der linearen Hülle $\mathcal{L}_x = \mathcal{L}(\mathbf{b}_1, \dots, \mathbf{b}_n)$ liegen. Sind $\mathcal{B}_1$ und $\mathcal{B}_2$ zwei Basen aus $\mathcal{L}_x$, so wird man sie wieder ineinander überführen können und es gelten die Beziehungen (3.250) und (3.251). Tatsächlich ist es so, dass man eine Basis für die Datenvektoren aus eben diesen Vektoren errechnen muß, – im allgemeinen hat man ja keine Information über die Basis außer der, die in den Daten steckt. Das bedeutet, dass man die Basisvektoren als Linearkombinationen der Datenvektoren berechnen muß, wobei zu berücksichtigen sein wird, dass die Anzahl r der benötigten Basisvektoren kleiner als n sein kann.

Analoge Betrachtungen gelten für die m n -dimensionalen Zeilenvektoren der Matrix X . Die Vektoren in einer Basis sind linear unabhängig, aber nicht notwendig orthogonal. Im Prinzip kann man irgendeine Basis zur "Erklärung" eines Datensatzes wählen, allerdings hat die Wahl einer orthogonalen Basis zumindest den Vorteil, dass man annehmen kann, dass die zu den Vektoren der Basis korrespondierenden Merkmale unkorreliert sind und damit unabhängig voneinander interpretiert werden können. Deswegen wird zunächst nur die Bestimmung einer orthogonalen Basis besprochen, von der man dann, wenn es gewünscht wird, zu einer nicht-orthogonalen Basis übergehen kann.

3.12 Bestimmung einer Basis für eine Datenmatrix

Es sei $X = [\mathbf{x}_1, \dots, \mathbf{x}_n]$ eine beliebige $(m \times n)$ Matrix mit dem Rang $r \leq \min(m, n)$. Gesucht ist eine Basis einerseits für die Spaltenvektoren von X , andererseits für die Zeilenvektoren von X . Nach Satz 3.2, 78, existieren stets Matrizen (m, r) - und (r, n) -Matrizen U und V ($U \in \mathbb{R}^{(m,r)}, V \in \mathbb{R}^{(r,n)}$) derart, dass $X = UV$. Die Spaltenvektoren von X sind Linearkombinationen der m -dimensionalen Spaltenvektoren $\mathbf{u}_k$ von U ($U = [\mathbf{u}_1, \dots, \mathbf{u}_r]$), und die Zeilenvektoren von X sind Linearkombinationen der Zeilenvektoren $\mathbf{v}_k$ von V ($V = [\mathbf{v}_1, \dots, \mathbf{v}_r]$), d.h. die Spalten von U sind Basisvektoren für die Spalten von X und die Zeilenvektoren von V sind Basisvektoren für die Zeilen von X . Es zeigt sich, dass die Wahl einer bestimmten Basis U die Wahl von V festlegt und umgekehrt.

Es gibt beliebig viele Möglichkeiten für eine Wahl von U bzw. V , so dass die Entscheidung für die Wahl einer bestimmten Basis nach irgendwelchen Kriterien erfolgen muß. So sei etwa X so definiert, dass die Zeilen von X zu "Fällen"²³ korrespondieren und die Spalten von X zu "Variablen"; für jeden Fall i ($i = 1, \dots, m$) gibt es dann n Messungen $x_{i1}, \dots, x_{in}$; die Koordinatenachsen repräsentieren die Variablen $V_1, \dots, V_n$. Der Zeilenvektor $(x_{i1}, \dots, x_{in})$ korrespondiert zum i -ten Zei-

²³Ein Fall ist eine Person oder ein Objekt, an dem Messungen vorgenommen werden, oder ein Zeitpunkt, zu dem die Messungen durchgeführt werden.

lenvektor $(u_{i1}, \dots, u_{ir})$ von U ; es gilt

$$(x_{i1}, \dots, x_{in}) = (u_{i1}, \dots, u_{ir})V \quad (3.253)$$

Die erste Spalte von U enthält dann die Koordinaten $u_{11}, \dots, u_{m1}$ der Fälle auf der ersten Achse eines durch die Basisvektoren U definierten Koordinatensystems, der zweite Spaltenvektor von U enthält die Koordinaten $u_{12}, \dots, u_{m2}$ der Fälle auf der zweiten Achse des durch Basisvektoren von U definierten Koordinatensystems, etc.

Die Frage ist nun, welchen Annahmen bezüglich U und V gemacht werden können, ohne die Allgemeinheit unnötig einzuschränken. So kann man postulieren, dass V orthonormal ist. Die $\mathbf{v}_1, \dots, \mathbf{v}_r$ bilden dann eine orthonormale Basis für die Zeilenvektoren von X ; von dieser Basis kann man, wenn es von Vorteil sein sollte, zu irgendeiner anderen, insbesondere auch nicht orthogonalen Basis übergehen. Jedenfalls folgt dann für eine orthonormale Matrix V aus $X = UV$ die Beziehung (Multiplikation von links mit V')

$$XV' = U, \quad (3.254)$$

und weiter wegen $U' = VX'$

$$VX'XV' = U'U. \quad (3.255)$$

Für den ersten Zeilenvektor $\mathbf{v}'_1$ von V gilt dann

$$\mathbf{v}'_1(X'X)\mathbf{v}_1 = \mathbf{u}'_1\mathbf{u}_1 \in \mathbb{R}. \quad (3.256)$$

Allerdings liegt damit noch nicht die Orientierung von $\mathbf{v}_1$ fest. Eine sinnvolle Forderung für die Orientierung ist, sie so zu wählen, dass $\mathbf{u}'_1\mathbf{u}_1$ maximal im Sinne des Satzes von Courant-Fischer ist (die Bedeutung dieser Annahme wird weiter unten noch elaboriert). Dann folgt, dass $\mathbf{v}_1$ gleich dem ersten Eigenvektor $\mathbf{t}_1$ von $X'X$ ist, und $\mathbf{u}'_1\mathbf{u}_1 = \lambda_1$ der zugehörige Eigenwert ist.

Ist T die orthonormale Matrix der Eigenvektoren von $X'X$ und Λ die Diagonalmatrix der zugehörigen Eigenwerte von $X'X$, so kann man $V = T'$ setzen. Dann folgt wegen $X'XT = T\Lambda$ und damit $X'X = T\Lambda T'$ aus $X = UT'$ die Beziehung $XT = U$ und

$$U'U = T'X'XT = T'(T\Lambda T')T = \Lambda, \quad (3.257)$$

d.h. die Matrix U ist orthogonal, da ja Λ eine Diagonalmatrix ist. Hat man die Eigenvektoren T und die zugehörigen Eigenwerte Λ bestimmt, so liegt mit (3.254) und (3.255) auch U als orthogonale Matrix fest. T und Λ müssen numerisch bestimmt werden; die Statistikpakete enthalten entsprechende Programme. Die Lösung $U = XT$ für U ist für T spezifisch, weshalb im Folgenden eine spezifische Bezeichnung gewählt wird: $U = L$, wobei L an den Ausdruck "latente Variable" erinnern soll. Dementsprechend hat man

$$X = LT'. \quad (3.258)$$

T definiert eine Hauptachsentransformation.

Alternativ dazu kann man fordern, dass die $\lambda_1, \dots, \lambda_n$ maximal im Sinne des Satzes von Courant-Fischer sein sollen. Dann folgt ebenfalls, dass T die Matrix der Eigenvektoren von $X'X$ ist.

Es sei X spaltenzentriert. $\vec{1}$ sei der m -dimensionale Einsvektor; dann ist $\vec{1}'\mathbf{x}_j = 0$ für alle Spaltenvektoren von X , oder $\vec{1}'X = \vec{0}' = (0, \dots, 0)$, $\vec{0}$ der n -dimensionale Nullvektor. Aus $XT = L$ folgt dann

$$\vec{1}XT = (\vec{1}'X)T = \vec{1}'L = \vec{0}', \quad (3.259)$$

d.h. die Spaltensummen von L sind ebenfalls alle gleich Null. Die Komponenten von $\mathbf{L}_k$ sind die Koordinaten der Fälle auf der k -ten latenten Dimension. Dann ist $\lambda_k = \mathbf{L}_k'\mathbf{L}_k$ proportional zur Varianz der Komponenten des k -ten Vektors $\mathbf{L}_k$, d.h. proportional zur Varianz der Koordinaten der Fälle auf dieser Dimension..

Ellipsoide für Datenpunkte: $\tilde{\mathbf{x}}_i = (x_{i1}, \dots, x_{in})'$ definiert einen Punkt im n -dimensionalen Variablenraum, der den i -ten Fall repräsentiert. Dem Punkt entspricht ein Ellipsoid, auf dem der Punkt liegt. Weiter sei $\tilde{\mathbf{L}}_i = (L_{i1}, \dots, L_{in})'$; man bemerke, dass $\tilde{\mathbf{x}}_i$ und $\tilde{\mathbf{L}}_i$ als *Spaltenvektoren* definiert worden sind. Aus (3.253) und $XT = L$, folgt (mit $V = T'$)

$$T'\tilde{\mathbf{x}}_i = \tilde{\mathbf{L}}_i. \quad (3.260)$$

Multiplikation von links mit $\tilde{\mathbf{L}}_i'\Lambda$ ergibt $\tilde{\mathbf{L}}_i'\Lambda T'\tilde{\mathbf{x}}_i = \tilde{\mathbf{L}}_i'\Lambda\tilde{\mathbf{L}}_i$, d.h.

$$\tilde{\mathbf{x}}_i'T\Lambda T'\tilde{\mathbf{x}}_i = \tilde{\mathbf{L}}_i'\Lambda\tilde{\mathbf{L}}_i, \quad (3.261)$$

und da $T\Lambda T' = X'X$ folgt

$$\tilde{\mathbf{x}}_i'(X'X)\tilde{\mathbf{x}}_i = \tilde{\mathbf{L}}_i'\Lambda\tilde{\mathbf{L}}_i = k_i. \quad (3.262)$$

k_i ist eine für den i -ten Fall charakteristische Konstante. Dementsprechend definiert $\{\mathbf{x}|\mathbf{x}'(X'X)\mathbf{x} = k_i\}$ ein Ellipsoid, auf dem der Punkt $\tilde{\mathbf{x}}_i$ liegt, und $\{\tilde{\mathbf{L}}|\tilde{\mathbf{L}}'\Lambda\tilde{\mathbf{L}} = k_i\}$ definiert ein achsenparalleles Ellipsoid, auf dem der zu $\tilde{\mathbf{x}}_i$ korrespondierende Punkt $\tilde{\mathbf{L}}_i$ liegt. Abbildung 9 zeigt schematisch die zu verschiedenen Punkten korrespondierenden Ellipsen mit identischer Orientierung. Abbildung 10 zeigt eine Punktekonfiguration aus einer 2-dimensional normalverteilten Population; in der linken Abbildung sind die Regressionsgeraden eingezeichnet, die von den Hauptachsen der Ellipsen, die durch die Kovarianzmatrix bestimmt werden, deutlich abweichen. Rechts wird die auf Achsenparallelität rotierte Konfiguration gezeigt.

Anmerkung zur Punktekonfiguration: Es ist gezeigt worden, dass für jeden Punkt der Konfiguration der Fälle ein Ellipsoid existiert, auf dem der Punkt liegt, aber dies bedeutet *nicht*, dass die Konfiguration auch tatsächlich ellipsoid sein muß, – d.h. die unterliegende Verteilung muß nicht die multivariate Normalverteilung sein. Auch wenn die Konfiguration nicht ellipsoid ist kann stets eine

Abbildung 9: Punktekongfiguration und zugehörige Ellipsen

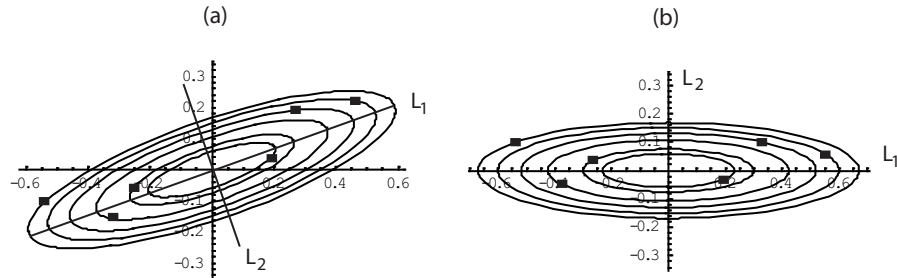
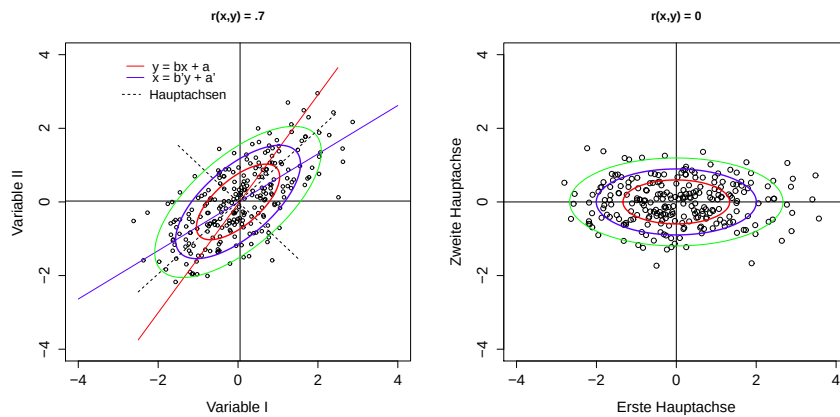


Abbildung 10: Links: Punktekongfiguration für $r_{xy} = .7$ mit Regressionsgeraden, Ellipsen und deren Hauptachsen; rechts: Die Hauptachsen als neue Koordinaten für die Punktekongfiguration.

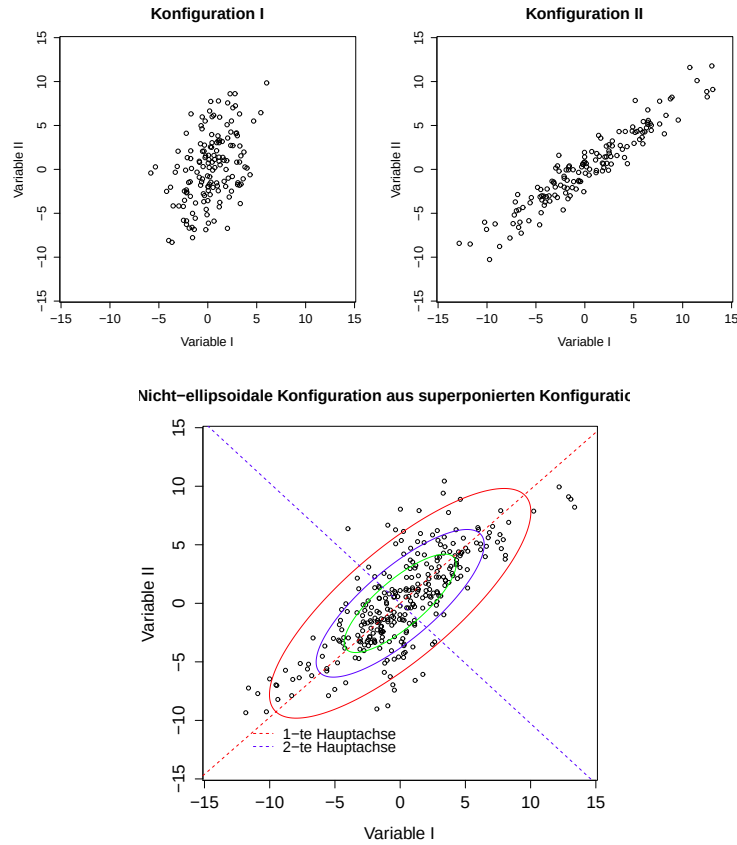


Menge von Ellipsoiden gefunden derart, dass jeder Fall auf einem Ellipsoid liegt, – einfach weil $X'X$ stets eine Menge von Ellipsoiden definiert. Dieser Fall kann eintreten, wenn sich die Stichprobe der Fälle aus Stichproben aus verschiedenen Populationen zusammensetzt. In Abbildung 11 wird dieser Sachverhalt illustriert.

3.13 Singularwertzerlegung und PCA

Gegeben sei eine $(m \times n)$ -Matrix X , deren Elemente x_{ij} Messwerte an $i = 1, \dots, m$ Fällen von $j = 1, \dots, n$ Variablen seien. Die Messungen der Variablen sind im Allgemeinen korreliert und man sucht eine möglichst geringe Anzahl von "latenten Variablen", mit denen die Korrelationen oder Kovarianzen "erklärt" werden kön-

Abbildung 11: Superponierte Punktekongfigurationen und Ellipsen



nen. Die latenten Variablen werden durch Vektoren repräsentiert, aus denen sich die Spalten- bzw. Zeilenvektoren von X als Linearkombinationen ergeben. Es sei $n < m$. Von einem formalen Standpunkt aus betrachtet man die m -dimensionalen Spaltenvektoren von X als Elemente eines $\mathbb{R}^m$, die sich als Linearkombinationen einer Teilbasis des $\mathbb{R}^m$ ergeben, – da es nur $n < m$ Spaltenvektoren gibt, genügt stets eine Teilbasis, weil n Vektoren notwendigerweise in einem maximal n -dimensionalen Teilraum des $\mathbb{R}^m$ liegen.

Die Spaltenvektoren $\mathbf{x}_1, \dots, \mathbf{x}_n$ von X repräsentieren ein Koordinatensystem, dass durch die gemessenen Variablen $\mathcal{V}_1, \dots, \mathcal{V}_n$ definiert ist: die Komponente x_{ij} ist die Koordinate des i -ten Falls auf der durch $\mathcal{V}_j$ definierten Achse. Gesucht ist eine Transformation dieser Koordinaten derart, dass (i) die neuen Achsen orthogonal sind, (ii) die Projektion der Datenpunkte $(x_{i1}, x_{i2}, \dots, x_{in})$ auf die erste neue Achse eine maximale Varianz hat, die auf der zweiten neuen Achse die zweitgrößte Varianz hat, etc. Damit ist aber auch schon klar, welche Repräsentation

von X diesen Forderungen gerecht wird: die SVD

$$X = Q\Sigma P', \quad \Sigma = \Lambda^{1/2}, \quad (3.263)$$

Man kann diese Gleichung in zwei Varianten anschreiben:

$$X = QA', \quad A = P\Sigma \quad (3.264)$$

$$= LP', \quad L = Q\Sigma \quad (3.265)$$

Es werde zuerst (3.264) betrachtet. Es folgt $Q'X = A'$, also

$$A = X'Q. \quad (3.266)$$

Es sei $\mathbf{p}_k$ der k -te Spaltenvektor von P ; wegen $X'X\mathbf{p}_k = \lambda_k\mathbf{p}_k$ repräsentiert $\mathbf{p}_k$ die k -te Hauptachse des zu $X'X$ korrespondierenden Ellipsoids, liefert also die Orientierung der k -ten "latenten" Dimension. Der Spaltenvektor $\mathbf{a}_k = (a_{1k}, a_{2k}, \dots, a_{nk})'$ enthält die Koordinaten der Variablen auf der k -ten latenten Dimension. Man spricht auch von "Ladungen" der Variablen auf der k -ten Dimension. Es gilt nun

$$A'A = Q'XX'Q' = \Lambda = \text{diag}(\lambda_1, \dots, \lambda_n), \quad n < m \quad (3.267)$$

d.h. es werden nur die zu Eigenwerten ungleich Null korrespondierenden Eigenvektoren von XX' betrachtet. Es seien $\mathbf{a}_k$ und $\mathbf{a}_s$ Spaltenvektoren von A . Es folgt $\mathbf{a}'_j\mathbf{a}_k = 0$ für $j \neq k$ und

$$\mathbf{a}'_k\mathbf{a}_s = \begin{cases} 0, & k \neq s \\ \|\mathbf{a}_k\|^2 = \lambda_k, & k = s \end{cases} \quad (3.268)$$

Man beachte, dass hier für eine gegebene Dimension über alle *Variablen* addiert wird.

Weiter gilt

$$AA' = X'QQ'X = X'X, \quad (3.269)$$

wegen der Orthonormalität von Q . Die Matrix X sei spaltenzentriert, so dass $\bar{\mathbf{1}}'X = \bar{\mathbf{0}}'$, d.h. die Spaltensummen von X sind gleich Null. Es seien $\tilde{\mathbf{a}}_j$, $\tilde{\mathbf{a}}_k$ der j -te und der k -te Zeilenvektor von A . Dann folgt

$$\frac{1}{m}\tilde{\mathbf{a}}'_j\tilde{\mathbf{a}}_k = \frac{1}{m}\mathbf{x}'_j\mathbf{x}_k = \begin{cases} c_{jk}, & j \neq k \\ s_j^2, & j = k \end{cases} \quad (3.270)$$

c_{jk} die Kovarianz der Variablen $\mathcal{V}_j$ und $\mathcal{V}_k$, s_j^2 die Varianz der j -ten Variablen. Man beachte, dass hier bei der Bildung des Skalarprodukts $\tilde{\mathbf{a}}'_j\tilde{\mathbf{a}}_k$ für gegebene j, k über die Dimensionen addiert wird. Ist X spaltenstandardisiert, so dass $\mathbf{x}_j = \mathbf{z}_j$, so impliziert (3.270)

$$\frac{1}{m}\tilde{\mathbf{a}}'_j\tilde{\mathbf{a}}_k = \begin{cases} r_{jk}, & j \neq k \\ 1, & j = k \end{cases} \quad (3.271)$$

Man kann die Variablen durch die Vektoren $\tilde{\mathbf{a}}_j$ graphisch repräsentieren. Die Repräsentation erlaubt u.U. eine erste Abschätzung der erforderlichen Anzahl r von latenten Dimensionen: geht man von standardisierten Messwerten aus, haben die $\tilde{\mathbf{a}}_j$ nach (3.271) alle die Länge $\|\tilde{\mathbf{a}}_j\| = 1$. Ist $r = 2$, so liegen die Endpunkte der $\tilde{\mathbf{a}}_j$ alle auf einem Kreis. Werden mehr als 2 Dimensionen benötigt, so liegen zumindest einige der Endpunkte innerhalb des Kreises, weil sie mindestens in eine dritte Dimension zeigen. Die Komponenten von $\mathbf{a}_k$ sind die Koordinaten der n Variablen auf der k -ten Dimension. Wegen $\lambda_1 \geq \lambda_2, \dots \geq \lambda_n$ ist dann $\|\mathbf{a}_1\|^2$ maximal, d.h. die Summe der quadrierten Koordinaten auf der ersten Dimension ist maximal, etc., für die folgenden Dimensionen wird die Summe stets kleiner, d.h. das Merkmal, das durch die erste Dimension repräsentiert wird, ist das bedeutendste, etc.

Es werde nun die Variante (3.265) betrachtet. Es ist $X'X = PL'LP' = P\Lambda P'$, d.h.

$$\mathbf{L}'_k \mathbf{L}_k = \|\mathbf{L}_k\|^2 = \lambda_k. \quad (3.272)$$

Die Komponenten von $\mathbf{L}_k$ sind die Koordinaten der Fälle auf der k -ten Dimension, und für zentrierte Matrix X ist $\bar{\mathbf{1}}'X = \bar{\mathbf{1}}'LP' = \bar{\mathbf{0}}'$, d.h. $\bar{\mathbf{1}}'L = \bar{\mathbf{0}}'$. Dann ist

$$\frac{1}{m} \|\mathbf{L}_k\|^2 = \tilde{s}_k^2 = \frac{1}{m} \lambda_k, \quad (3.273)$$

$\tilde{s}_k^2$ die Varianz der Koordinaten der Fälle auf der k -ten Dimension. Nach (3.268) gilt aber auch $\|\mathbf{a}_k\|^2 = \lambda_k$, und $\|\mathbf{a}_k\|^2$ ist die Summe der Quadrate der Koordinaten der Variablen auf der k -ten Dimension, d.h. es ist

$$m\tilde{s}_k^2 = \|\mathbf{a}_k\|^2 = \lambda_k. \quad (3.274)$$

Der Eigenwert λ_k ist also proportional zur Varianz der Koordinaten der Fälle auf der k -ten Dimension und gleich der Summe der Quadrate der Koordinaten der Variablen auf der k -ten Dimension (die Summe der Komponenten von $\mathbf{a}_k$ ist nicht notwendig gleich Null, daher kann bei $\|\mathbf{a}_k\|^2$ nicht von einer Varianz gesprochen werden).

Wie Korrelationen zwischen Variablen über die Elemente von A ausgedrückt werden können (vergl. (3.271)), so können Skalarprodukte zwischen Fällen über die Elemente von L angeschrieben werden: Es ist

$$LL' = Q\Sigma\Sigma'Q' = Q\Lambda Q' = XX',$$

so dass

$$\tilde{L}_i \tilde{L}'_j = \tilde{\mathbf{x}}_i \tilde{\mathbf{x}}'_j,$$

wobei $\tilde{L}_i, \tilde{L}_j$ Zeilenvektoren von L und $\tilde{\mathbf{x}}_i$ und $\tilde{\mathbf{x}}_j$ Zeilenvektoren von X sind. Damit $\tilde{\mathbf{x}}_i \tilde{\mathbf{x}}'_j$ aber eine Korrelation zwischen den Fällen i und j ist, muß X nicht spalten-, sondern zeilenstandardisiert sein. Darüber hinaus muß die Summation über die Variablen, die diese Skalarproduktbildung beinhaltet, sinnvoll sein.

(3.263) impliziert

$$X = \sum_{k=1}^n \sigma_j \mathbf{q}_k \mathbf{p}'_k, \quad \sigma_j = \sqrt{\lambda_j}. \quad (3.275)$$

Nach dem Satz von Eckart & Young (1936) liefert die SVD eine Approximation $X \approx Q_r \Sigma_r P'_r$ im Sinne der Kleinsten Quadrate, bei der nur die ersten r Eigenwerte und zugehörigen Eigenvektoren eingehen. Bildet man den Quotienten

$$\pi_k = \frac{\lambda_k}{\sum_{j=1}^n \lambda_j} \quad (3.276)$$

so gibt π_k wegen (3.274) den Anteil der Varianz an der Gesamtvarianz, der durch $\mathbf{f}_k$ erklärt wird. Ein Kriterium für die Wahl von r ist, denjenigen Wert von r zu wählen, für den $\sigma_{r+1} \ll \sigma_r$ ist (Scree-Test). Die Summe

$$E = \sum_{k=r+1}^n \sigma_j \mathbf{q}_k \mathbf{p}'_k \quad (3.277)$$

(vergl. (3.275)) wird dann als "Fehler" interpretiert.

Die Repräsentation $X = LP'$ eignet sich, wenn in erster Linie die Fälle graphisch repräsentiert werden sollen. In diesem Fall ist die Varianz der Koordinaten der Fälle auf der ersten Dimension maximal, etc. Man kann versuchen "Typen" zu identifizieren.

3.14 Eigenvektorberechnung und Deflation einer Matrix

Es sei A eine symmetrische $(n \times n)$ -Matrix; die Eigenvektoren $\mathbf{P}_1, \dots, \mathbf{P}_n$ von A sind dann orthogonal bzw. nach Normierung orthonormal. Obwohl genügend Programme zur aktuellen Berechnung der Eigenwerte und -vektoren zur Verfügung stehen, soll hier kurz gezeigt werden, wie eine solche Berechnung im Prinzip durchgeführt werden kann. Dieses Prinzip spielt in einigen multivariaten Verfahren, z.B. im Partial Least Square (PLS-) Verfahren eine grundsätzliche Rolle.

Deflation einer Matrix Zunächst wird der Begriff der Deflation einer Matrix vorgestellt. Es sei λ_1 der größte Eigenwert von A und $\mathbf{P}_1$ der zugehörige Eigenvektor. Dann heißt

$$\tilde{A} = A - \lambda_1 \mathbf{P}_1 \mathbf{P}'_1 \quad (3.278)$$

eine durch Deflation erzeugte Matrix, wobei $\mathbf{P}_1 \mathbf{P}'_1$ das dyadische Produkt von $\mathbf{P}_1$ mit sich selbst ist. Die $\mathbf{P}_j$ seien normiert. Es sei $\mathbf{P}_j$ ein weiterer Eigenvektor von A . Dann folgt

$$\tilde{A} \mathbf{P}_j = (A - \lambda_1 \mathbf{P}_1 \mathbf{P}'_1) \mathbf{P}_j = A \mathbf{P}_j - \lambda_1 \mathbf{P}_1 \mathbf{P}'_1 \mathbf{P}_j = \begin{cases} 0, & j = 1, \\ \lambda_j \mathbf{P}_j, & j \neq 1 \end{cases} \quad (3.279)$$

denn $\mathbf{P}_1' \mathbf{P}_1 = 1$ und $\mathbf{P}_1' \mathbf{P}_j = 0$ für $j \neq 1$. Der größte Eigenwert von $\tilde{A}$ ist als der zweitgrößte Eigenwert von A . Man kann dann $\tilde{A}$ noch einmal deflationieren, und der größte Eigenwert dieser deflationierten Matrix ist der drittgrößte Eigenwert von A , etc. Hat man eine Methode, den jeweils größten Eigenwert einer Matrix zu berechnen, so liefert die sukzessive Deflation der Matrizen nacheinander die Eigenwerte und zugehörigen Eigenvektoren einer Matrix.

Iteration: Es sei $\mathbf{z}_0$ ein beliebiger n -dimensionaler Vektor; dann ist $\mathbf{z}_1 = A\mathbf{z}_0$ ebenfalls ein n -dimensionaler Vektor. Man kann dann den Vektor $\mathbf{z}_2 = A\mathbf{z}_1 = A^2\mathbf{z}_0$ etc berechnen und hat für beliebiges k $\mathbf{z}_k = A^k\mathbf{z}_0$. Der Vektor $\mathbf{z}_0$ kann als Linearkombination der Eigenvektoren $\mathbf{P}_1, \dots, \mathbf{P}_n$ dargestellt werden, da die Eigenvektoren von A ja eine Basis für den $\mathbb{R}^n$ bilden; man hat etwa

$$\mathbf{z}_0 = a_1\mathbf{P}_1 + a_2\mathbf{P}_2 + \dots + a_n\mathbf{P}_n.$$

Dann folgt

$$\mathbf{z}_1 = A\mathbf{z}_0 = a_1A\mathbf{P}_1 + a_2A\mathbf{P}_2 + \dots + a_nA\mathbf{P}_n = a_1\lambda_1\mathbf{P}_1 + a_2\lambda_2\mathbf{P}_2 + \dots + a_n\lambda_n\mathbf{P}_n,$$

und

$$\begin{aligned} \mathbf{z}_2 &= a_1\lambda_1A\mathbf{P}_1 + a_2\lambda_2A\mathbf{P}_2 + \dots + a_n\lambda_nA\mathbf{P}_n \\ &= a_1\lambda_1^2\mathbf{P}_1 + a_2\lambda_2^2\mathbf{P}_2 + \dots + a_n\lambda_n^2\mathbf{P}_n, \end{aligned}$$

und schließlich

$$\mathbf{z}_k = A^k\mathbf{z}_0 = a_1\lambda_1^k\mathbf{P}_1 + a_2\lambda_2^k\mathbf{P}_2 + \dots + a_n\lambda_n^k\mathbf{P}_n.$$

Aber diese Gleichung kann in der Form

$$\mathbf{z}_k = \lambda_1^k \left(a_1\mathbf{P}_1 + \left(\frac{\lambda_2}{\lambda_1}\right)^k + \dots + \left(\frac{\lambda_n}{\lambda_1}\right)^k \mathbf{P}_n \right)$$

geschrieben werden. Da $\lambda_1 > \lambda_j$ für $j \neq 1$ folgt $(\lambda_j/\lambda_1)^k \rightarrow 0$ für größer werdendes k , so dass

$$\lim_{k \rightarrow \infty} \mathbf{z}_k \rightarrow a_1\lambda_1^k\mathbf{P}_1, \quad (3.280)$$

Man bestimmt nun den Rayleigh-Quotienten für $\mathbf{z}_k$:

$$\lim_{k \rightarrow \infty} \frac{\mathbf{z}_k' A \mathbf{z}_k}{\|\mathbf{z}_k\|^2} = \frac{a_1^2 \lambda_1^{2k} \mathbf{P}_1' A \mathbf{P}_1}{a_1^2 \lambda_1^{2k}} = \frac{a_1^2 \lambda_1^{2k} \mathbf{P}_1' \lambda_1 \mathbf{P}_1}{a_1^2 \lambda_1^{2k}} = \lambda_1 \quad (3.281)$$

wegen $\mathbf{P}_1' \mathbf{P}_1 = 1$. Mit $q_k = \mathbf{z}_k' A \mathbf{z}_k / \|\mathbf{z}_k\|^2$ erhält man mit

$$\lim_{k \rightarrow \infty} \frac{1}{q_k} \mathbf{z}_k = a_1 \mathbf{P}_1 \quad (3.282)$$

eine Abschätzung des ersten Eigenvektors $a_1\mathbf{P}_1$ (und durch Normierung von $\mathbf{P}_1$). In der gleichen Weise wendet man die Iteration auf die deflationierte Matrix $\tilde{A}$ an und erhält dann eine Abschätzung von λ_2 und $\mathbf{P}_2$, etc.

3.15 Die verallgemeinerte Inverse

Die Inverse A^{-1} einer Matrix ist nur definiert, wenn (i) A quadratisch ist, und (ii) vollen Rang hat, d.h. wenn A eine $n \times n$ -Matrix ist und der Rang r gleich n ist. Es ist aber manchmal von Interesse, die Inverse einer Matrix berechnen zu können, wenn entweder (i) oder (ii) oder weder (i) noch (ii) erfüllt sind.

Definition 3.21 *Es sei A eine $m \times n$ -Matrix und B eine $n \times m$ -Matrix. B heißt verallgemeinerte Inverse (oder Pseudoinverse), wenn*

$$ABA = A, \quad BAB = B \quad (3.283)$$

gilt.

Statt des Ausdrucks 'verallgemeinerte Inverse' ist auch der Ausdruck 'Pseudoinverse' oder 'Moore-Penrose-Inverse' üblich (Moore 1920), Penrose (1954)). Für B wird auch A^+ oder A^- geschrieben.

Beispiele:

1. Es sei $A\mathbf{x} = \mathbf{y}$ ein lineares Gleichungssystem, wobei A eine $m \times n$ -Matrix, und $\mathbf{y}$ eine Linearkombination der Spaltenvektoren von A sei; $\mathbf{x}$ sei unbekannt. Es sei $m \geq n = r$, d.h. A habe "vollen" Rang. Dann hat $A'A$ ebenfalls vollen Rang und die Inverse $(A'A)^{-1}$ zu $A'A$ existiert. Dann erhält man durch Multiplikation von links mit A' die Gleichung $A'A\mathbf{x} = A'\mathbf{y}$ und schließlich durch Multiplikation von links mit $(A'A)^{-1}$

$$\mathbf{x} = (A'A)^{-1}A'\mathbf{y} \Rightarrow A^- = (A'A)^{-1}A'. \quad (3.284)$$

Dass A^- tatsächlich eine generalisierte Inverse ist, sieht man durch einsetzen:

$$A((A'A)^{-1}A')A = A(A'A)^{-1}A'A = A,$$

d.h. die Bedingung (3.283) ist erfüllt.

2. Es sei $A = Q\Sigma P'$, $\Sigma = \Lambda^{1/2}$; $Q\Sigma P'$ ist die SVD von A . Dann ist

$$(A'A)^{-1}A' = (P\Lambda P')^{-1}P\Lambda^{1/2}Q' = P\Lambda^{-1}P'P\Lambda^{1/2}Q' = P\Lambda^{-1/2}Q',$$

d.h.

$$A^- = P\Lambda^{-1/2}Q' = P\Sigma^{-1}Q' \quad (3.285)$$

ist eine generalisierte Inverse.

3. Die Anwendung der SVD erlaubt die Definition einer generalisierten Inversen für den Fall $r < \min(m, n)$. Dann gilt

$$A = Q_r \Sigma_r P_r'$$

wobei Q_r nur die ersten r Eigenvektoren von $A'A$ enthält, P_r enthält nur die ersten r Eigenvektoren von $P'P$ und $\Sigma_r = \text{diag}(\sigma_1, \dots, \sigma_r)$, mit $\sigma_j = \sqrt{\lambda_j}$ und λ_j der j -te von Null verschiedene Eigenwert von AA' bzw. $A'A$. Analog zu (3.285) findet man

$$A^- = Q_r \Sigma_r^{-1} P_r'. \quad (3.286)$$

3.16 Lineare Gleichungssysteme

3.16.1 Allgemeine Charakterisierung der Lösungen

Es sei A eine $(m \times n)$ -Matrix, $\mathbf{x}$ ein n -dimensionaler Vektor, und $\mathbf{y}$ ein m -dimensionaler Vektor, und es gelte

$$A\mathbf{x} = \mathbf{y}, \quad \mathbf{x} \in \mathbb{R}^n, \quad \mathbf{y} \in \mathbb{R}^m \quad (3.287)$$

Diese Gleichung kann als ein System von m linearen Gleichungen mit n Unbekannten, nämlich den Komponenten von $\mathbf{x}$ gesehen werden. $\mathbf{x}$ ist n -dimensional, $\mathbf{y}$ ist m -dimensional. Sind A und $\mathbf{y}$ vorgegeben, so stellt sich die Frage, ob es überhaupt einen Lösungsvektor $\mathbf{x}$ gibt, und wenn ja, ob es mehrere Lösungsvektoren gibt. Man kann zwischen zwei Arten von Gleichungssystemen unterscheiden:

1. $\mathbf{y} = \vec{0}$. Das Gleichungssystem heißt dann *homogen*,
2. $\mathbf{y} \neq \vec{0}$. Das Gleichungssystem heißt dann *inhomogen*.

Definition 3.22 Es sei A eine $(m \times n)$ -Matrix mit dem Rang $r = \text{rg}(A) \leq \min(m, n)$. und $A\mathbf{x} = \mathbf{y}$ sei ein System von Gleichungen. Weiter sei

$$\text{kern}(A) = \{\mathbf{x} \in \mathbb{R}^n | A\mathbf{x} = \vec{0}\} \quad (3.288)$$

$$\mathcal{L}(A) = \{\mathbf{y} \in \mathbb{R}^m | A\mathbf{x} = \mathbf{y}\}. \quad (3.289)$$

$\text{kern}(A)$ heißt Kern von A , und $\mathcal{L}(A)$ ist die lineare Hülle von A .

Anmerkungen:

1. In Definition 3.2, Punkte 4., Seite 71 ist der Begriff des Kerns einer Abbildung eingeführt worden. Die Matrix A definiert eine Abbildung, und (3.288) definiert damit den Kern einer Abbildung. $\text{kern}(A)$ ist ein (Teil-)Vektorraum: sind die Vektoren $\mathbf{x}_1$ und $\mathbf{x}_2$ Elemente aus $\text{kern}(A)$, so rechnet man leicht nach, dass dann auch $a_1\mathbf{x}_1 + a_2\mathbf{x}_2 \in \text{kern}(A)$ ($a_1, a_2 \in \mathbb{R}$) gilt.
2. Die Gleichung $A\mathbf{x} = \mathbf{y} \neq \vec{0}$ bedeutet, dass $\mathbf{y}$ als Linearkombination der Spalten von A dargestellt werden soll: mit $\mathbf{x} = (x_1, \dots, x_n)'$ ist

$$A\mathbf{x} = x_1\mathbf{a}_1 + \dots + x_n\mathbf{a}_n = \mathbf{y}, \quad (3.290)$$

wobei $\mathbf{a}_1, \dots, \mathbf{a}_n$ die Spaltenvektoren von A sind; damit also (3.290) gilt (d.h. damit eine Lösung $\mathbf{x}$ existiert), müssen die Matrizen A und $(A, \mathbf{y})$ (die um die Spalte $\mathbf{y}$ erweiterte Matrix A) denselben Rang haben. $\square$

Der folgende Satz gilt für beliebige $(m \times n)$ -Matrizen X ; er wird hier für $X = A$ angeschrieben, weil er in Bezug auf das Gleichungssystem $A\mathbf{x} = \mathbf{y}$ interpretiert werden soll:

Satz 3.27 *Es sei A eine (m, n) -Matrix mit der SVD $A = Q\Sigma T'$, wobei Q aus den Spaltenvektoren $\mathbf{q}_1, \dots, \mathbf{q}_n$ und T aus den Spaltenvektoren $\mathbf{t}_1, \dots, \mathbf{t}_n$ bestehe; ist $\text{rg}(A) = r \leq \min(m, n)$, so sind r Singularwerte σ_k größer als Null und $n - r$ Singularwerte sind gleich Null. Dann gilt*

$$\mathcal{L}(A) = \mathcal{L}(\mathbf{q}_1, \dots, \mathbf{q}_r) \quad (3.291)$$

$$\text{kern}(A) = \mathcal{L}(\mathbf{t}_{r+1}, \dots, \mathbf{t}_n), \quad s = \min(m, n) \quad (3.292)$$

$$\text{rg}(\text{kern}(A) + \text{rg}(\mathcal{L}(A))) = \min(m, n) \quad (3.293)$$

Beweis: Der Beweis wird für den Fall $n \leq m$ (höchstens so viele Unbekannte wie Gleichungen) geführt; der Beweis für den Fall $m < n$ (mehr Unbekannte als Gleichungen) ist analog. Wegen $A = Q\Sigma T'$ sind die $\mathbf{a}_j$ Linearkombinationen der $r \leq \min(m, n)$ Spaltenvektoren $\mathbf{q}_k$ von Q ; als Eigenvektoren von AA' sind die $\mathbf{q}_k$ paarweise orthogonal und damit linear unabhängig; sie bilden eine r -dimensionale Teilbasis des $\mathbb{R}^m$. Damit sind auch alle Linearkombinationen der $\mathbf{a}_j$ als Linearkombinationen der $\mathbf{q}_k$ darstellbar, so dass (3.291) gelten muß.

Der Kern von A sind alle n -dimensionalen Vektoren $\mathbf{x}$, für die $A\mathbf{x} = \vec{0}$ gilt. Die Spaltenvektoren von T bilden eine Basis des $\mathbb{R}^n$, so dass allgemein

$$\mathbf{x} = c_1 \mathbf{t}_1 + \dots + c_r \mathbf{t}_r + c_{r+1} \mathbf{t}_{r+1} + \dots + c_n \mathbf{t}_n$$

geschrieben werden kann, und es gilt

$$A\mathbf{x} = A \sum_{j=1}^n c_j \mathbf{t}_j = \sum_{j=1}^n c_j A\mathbf{t}_j = \sum_{j=1}^r c_j \sigma_j \mathbf{q}_j = \vec{0},$$

denn $\sigma_j = 0$ für $j > r$ (falls $r < \min(m, n)$) (vergl. (3.176), Seite 118). Wegen der linearen Unabhängigkeit der $\mathbf{q}_j$ kann diese Gleichung nur gelten, wenn $c_1 = \dots = c_r = 0$. Dann kann $\mathbf{x} \neq \vec{0}$ kein Element des durch die $\mathbf{t}_1, \dots, \mathbf{t}_r$ aufgespannten Vektorraums sein, sondern muß ein Element des $(n - r)$ -dimensionalen Komplementärums sein. Die $\mathbf{t}_{r+1}, \dots, \mathbf{t}_n$ sind eine Basis für diesen Komplementärums, so dass man

$$\mathbf{x} = c_{r+1} \mathbf{t}_{r+1} + \dots + c_n \mathbf{t}_n,$$

ansetzen kann, und

$$A\mathbf{x} = A \sum_{j=r+1}^n c_j \mathbf{t}_j = \sum_{j=r+1}^n c_j A\mathbf{t}_j = \sum_{j=r+1}^n c_j \sigma_j \mathbf{q}_j = \vec{0},$$

wegen (3.176), Seite 118, und weil $\sigma_j = 0$ für $j > r$, falls $r < \min(m, n)$. Alle Linearkombinationen von Vektoren $\mathbf{x} \neq \vec{0}$ mit $A\mathbf{x} = \vec{0}$ sind Linearkombinationen der $\mathbf{t}_{r+1}, \dots, \mathbf{t}_n$, und dies ist die Aussage von (3.292).

Die Gleichung (3.293) ist eine unmittelbare Folge der vorangegangenen Argumente: $\mathcal{L}(A)$ hat den Rang r und $\text{kern}(A)$ hat den Rang $n - r$, so dass die Summe der Ränge gleich n sein muß. $\square$

Anmerkung: Der Satz 3.27 ergab sich als Folgerung aus der SVD für die Matrix A . Die Eigenvektoren $\mathbf{t}_j$ von $A'A$ sind aber nicht die einzigen Vektoren, mit denen sich $\text{kern}(A)$ darstellen läßt. Einen alternativen, wenn auch etwas länglichen alternativen Beweis, in dem ein anderer Satz von Basisvektoren verwendet wird, findet man im Anhang, Abschnitt 8.5. $\square$

Es bleibt noch, die allgemeine Lösungsmenge zu spezifizieren:

Satz 3.28 *Es sei $A\mathbf{x} = \mathbf{y} \in \mathcal{L}(A)$, $\text{rg}(A) = r$, und insbesondere sei $\mathbf{x} = \mathbf{x}_0$ eine bestimmte Lösung, so dass $A\mathbf{x}_0 = \mathbf{y}$ gilt. Der Kern $\text{kern}(A)$ besteht aus dem $(n - r)$ -dimensionalen Teilraum $\mathcal{L}_{n-r} = \mathcal{L}(\mathbf{t}_{r+1}, \dots, \mathbf{t}_n)$ des V_n . Dann ist die Menge der Lösungsvektoren durch*

$$\mathcal{L} = \{\mathbf{x}_0 + \mathbf{x} | \mathbf{x} \in \mathcal{L}_{n-r}\} \quad (3.294)$$

gegeben.

Beweis: Tatsächlich ist $\mathbf{x}_0 + \mathbf{x}$ eine Lösung, denn

$$A(\mathbf{x}_0 + \mathbf{x}) = A\mathbf{x}_0 + A\mathbf{x} = A\mathbf{x}_0,$$

denn $A\mathbf{x} = \vec{0}$ ist nach Voraussetzung eine Lösung, und $\mathbf{x}_0$ war als Lösungsvektor vorausgesetzt worden. Umgekehrt sei $\mathbf{x}_1$ ein Lösungsvektor. Es muß gezeigt werden, dass $\mathbf{x}_1 \in \mathcal{L}$ ist. Nach Voraussetzung muß $A\mathbf{x}_1 = \mathbf{y}$ gelten. Für irgendeinen Vektor $\mathbf{x} \in \mathcal{L}_{n-r}$ muß $A\mathbf{x} = \vec{0}$ gelten. Dann muß aber auch

$$A(\mathbf{x}_1 + \mathbf{x}) = A\mathbf{x}_1 = \mathbf{y}$$

gelten, so dass $\mathbf{x}_1 + \mathbf{x} \in \mathcal{L}$ liegt. $\square$

Der Fall $m = n = r$: Ist $m = n = r$, r der Rang von A , so existiert die Inverse A^{-1} und aus $A\mathbf{x} = \mathbf{y} \in V_n^r = \mathcal{L}(A)$ folgt sofort die Lösung

$$\mathbf{x} = A^{-1}\mathbf{y}, \quad \mathbf{y} \in \mathcal{L}(A). \quad (3.295)$$

Der Fall $m > n = r$: In diesem Fall gibt es mehr Gleichungen als Unbekannte; das Gleichungssystem ist überbestimmt. Im Allgemeinen wird man keinen Lösungsvektor $\mathbf{x}$ finden, der allen Gleichungen exakt genügt. Dies ist z.B. bei der multiplen Regression der Fall, da man üblicherweise eine größere Anzahl m von Fällen als unbekannte Regressionsparameter hat. Die $(m \times n)$ -Matrix $A = X$

der Prädiktoren hat aber im Allgemeinen den vollen Rang $r = n$, so dass man eine Lösung finden könnte, indem man von links mit A' multipliziert, so dass $A'Ax = A'y$ folgt, und da $A'A$ den gleichen Rang wie A hat (Satz ??, Seite ??) existiert die zu $A'A$ inverse Matrix $(A'A)^{-1}$, so dass

$$\mathbf{x} = (A'A)^{-1}A'y \quad (3.296)$$

resultiert. Sind die Komponenten von $\mathbf{y}$ Messwerte, so sind sie üblicherweise durch Messfehler kontaminiert, so dass $\mathbf{y} \notin \mathcal{L}(A)$. (3.296) liefert dann keine Lösung $\mathbf{x}$, die allen m Gleichungen genügt. (3.296) ist dann die bekannte Kleinste-Quadrate-Schätzung $\hat{\mathbf{x}}$ für $\mathbf{x}$.

Für den Fall $r < n$ liefert (3.294) den Lösungsraum. Zur Berechnung der Lösung s. a. Abschnitt 3.7.1, wo die *Cramersche Regel* eingeführt wird.

3.16.2 Die Cramersche Regel

In Abschnitt 3.7.1 wurde schon kurz die Cramersche Regel eingeführt: Nach (3.79), Seite 84, hat man für ein Gleichungssystem mit zwei Unbekannten die Lösungen

$$x_j = \frac{|A_j|}{|A|}, \quad j = 1, 2,$$

wobei die Matrix A_j aus der Matrix A entsteht, wenn man die j -te Zeile und die j -te Spalte streicht.

Die Beziehung kann für den Fall eines Systems mit n Unbekannten verallgemeinert werden. Das Gleichungssystem $A\mathbf{x} = \mathbf{y}$ besteht ja aus n Gleichungen

$$\sum_{j=1}^n a_{jk}x_j = y_j, \quad i = 1, \dots, n.$$

Es werde vorausgesetzt, dass A den Rang n hat. Es sei A_j die Matrix, die entsteht, wenn man die j -te Spalte von A durch $\mathbf{y}$ ersetzt. Dann ist

$$x_j = \frac{|A_j|}{|A|}, \quad j = 1, \dots, n \quad (3.297)$$

x_j die j -te Komponente von $\mathbf{x}$, also die j -te Unbekannte. Dies ist die *Cramersche Regel*²⁴.

Beispiel 3.16 Berechnung der inversen Matrix In Beispiel 3.4 wurde gezeigt, dass die Berechnung der Inversen A^{-1} einer quadratischen Matrix X auf das Lösen linearer Gleichungssysteme hinausläuft: da $AA^{-1} = I$ muß für den j -ten Spaltenvektor $\tilde{\mathbf{a}}_j$ von A^{-1} die Beziehung

$$A\tilde{\mathbf{a}}_j = \mathbf{e}_j \quad (3.298)$$

²⁴Gabriel Cramer (1702 – 1742), Schweizer Mathematiker

gelten, $\mathbf{e}_j$ der j -te Spaltenvektor von I , $\mathbf{e}_j = (0, \dots, 0, 1, 0, \dots, 0)'$. Nach der Cramerschen Regel erhält man dann für die i -te Komponente von $\tilde{\mathbf{a}}_j$

$$\tilde{a}_{ij} = \frac{|A_i|}{|A|} \quad (3.299)$$

Wird nun $|A_i|$ nach dem Laplaceschen Entwicklungssatz nach der i -ten Spalte entwickelt, so folgt

$$\tilde{a}_{ij} = \frac{(-1)^{i+j} |A_{ji}|}{|A|}, \quad (3.300)$$

d.h. die Elemente $\tilde{a}_{ij}$ von A^{-1} sind bis auf den Faktor $1/|A|$ gleich den korrespondierenden Kofaktoren von A . Dieser Sachverhalt erweist sich als wichtig für die Interpretation der Inversen von Korrelationsmatrizen $\square$

3.16.3 Lineare Gleichungen und Gauß-Algorithmus

Gegeben sei ein lineares Gleichungssystem $A\mathbf{x} = \mathbf{y}$; gesucht ist eine möglichst systematische Art und Weise, einen Lösungsvektor $\mathbf{x}$ zu bestimmen. In diesem Abschnitt soll das gaußsche Eliminationsverfahren (auch Gauß-Algorithmus) vorgestellt werden. Die Idee des Verfahrens ist, die Koeffizientenmatrix A in eine Dreiecksmatrix zu transformieren, anhand der sehr schnell eine Lösung gefunden werden kann, sofern eine Lösung existiert. Zur Illustration werde ein System mit drei Unbekannten betrachtet:

$$\begin{aligned} a_{11}x_1 + a_{12}x_2 + a_{13}x_3 &= y_1 \\ a_{21}x_1 + a_{22}x_2 + a_{23}x_3 &= y_2 \\ a_{31}x_1 + a_{32}x_2 + a_{33}x_3 &= y_3 \end{aligned} \quad (3.301)$$

Bei der Bestimmung der Dreiecksmatrix mit den Elementen $\tilde{a}_{ij}$ wird auch der Vektor $\mathbf{y}$ in einen Vektor $\tilde{\mathbf{y}}$ transformiert:

$$\begin{aligned} \tilde{a}_{11}x_1 + \tilde{a}_{12}x_2 + \tilde{a}_{13}x_3 &= \tilde{y}_1 \\ \tilde{a}_{22}x_2 + \tilde{a}_{23}x_3 &= \tilde{y}_2 \\ \tilde{a}_{33}x_3 &= \tilde{y}_3 \end{aligned} \quad (3.302)$$

Für dieses System findet man sofort die Lösung für x_3 , nämlich $x_3 = \tilde{y}_3/\tilde{a}_{33}$, die man dann in die zweite Gleichung einsetzen kann, die dann nur noch die Unbekannten x_1 und x_2 enthält und die damit nach x_2 aufgelöst werden kann, und die Lösungen für x_3 und x_2 können dann in die erste Gleichung eingesetzt werden, um x_1 zu bestimmen.

Um die Koeffizienten $\tilde{a}_{ij}$ zu bestimmen, macht man von den *elementaren Umformungen* Gebrauch. Demnach bleibt eine Gleichung ja korrekt, wenn man beide Seiten mit einem Faktor multipliziert. Ebenso kann man man zwei Gleichungen

addieren; es entsteht wieder eine gültige Gleichung. Diese Operationen muß man so anwenden, dass man von den a_{ij} zu den $\tilde{a}_{ij}$ gelangt.

Die Koeffizientenmatrix in (3.302) ist eine Dreiecksmatrix

$$\tilde{A} = \begin{pmatrix} \tilde{a}_{11} & \tilde{a}_{12} & \cdots & \tilde{a}_{1n} \\ 0 & \tilde{a}_{22} & \cdots & \tilde{a}_{2n} \\ & & \ddots & \\ 0 & 0 & \cdots & \tilde{a}_{nn} \end{pmatrix}. \quad (3.303)$$

Die elementaren Umformungen, die von A zu $\tilde{A}$ führen, verändern den Rang nicht, d.h. $\tilde{A}$ hat denselben Rang wie A . Berechnet man die Determinante von $\tilde{A} - \lambda I_n$, so erhält man

$$|\tilde{A}| = \begin{vmatrix} \tilde{a}_{11} - \lambda & \tilde{a}_{12} & \cdots & \tilde{a}_{1n} \\ 0 & \tilde{a}_{22} - \lambda & \cdots & \tilde{a}_{2n} \\ & & \ddots & \\ 0 & 0 & \cdots & \tilde{a}_{nn} - \lambda \end{vmatrix} = (\tilde{a}_{11} - \lambda)(\tilde{a}_{22} - \lambda) \cdots (\tilde{a}_{nn} - \lambda) = 0. \quad (3.304)$$

Daraus folgt unmittelbar, dass die $\tilde{a}_{jj}$ die Eigenwerte von $\tilde{A}$ sind.

Die Anwendung der elementaren Umformungen macht man sich am besten durch ein konkretes Beispiel klar:

$$\begin{aligned} 2x_1 - x_2 + 3x_3 &= 1 & (I) \\ 3x_1 + x_2 - 2x_3 &= 0 & (II) \\ x_1 + x_2 + x_3 &= 3 & (III) \end{aligned}$$

Die Koeffizientenmatrix ist

$$\left(\begin{array}{ccc|c} 2 & -1 & 3 & 1 \\ 3 & 1 & -2 & 0 \\ 1 & 1 & 1 & 3 \end{array} \right)$$

Man beginnt, indem man die erste Gleichung (I) noch einmal anschreibt, die zweite durch $3III - II$ ersetzt und die dritte durch $2III - I$

$$\begin{aligned} 2x_1 - x_2 + 3x_3 &= 1 \\ 0 + 2x_2 + 5x_3 &= 9 \\ 0 + 3x_2 + x_3 &= 5 \end{aligned}$$

Die Koeffizientenmatrix ist hier

$$\left(\begin{array}{ccc|c} 2 & -1 & 3 & 1 \\ 0 & 2 & 5 & 9 \\ 0 & 3 & 1 & 5 \end{array} \right)$$

Ersetzt man die Gleichung (3.305) durch $3(3.305) - 2(3.305)$, so erhält man

$$\begin{aligned} 2x_1 - x_2 + 3x_3 &= 1 \\ 0 + 2x_2 + 5x_3 &= 9 \\ 0 + 0 + 17x_3 &= 17 \end{aligned}$$

mit der Koeffizientenmatrix

$$\left(\begin{array}{ccc|c} 2 & -1 & 3 & 1 \\ 0 & 2 & 5 & 9 \\ 0 & 0 & 17 & 17 \end{array} \right) \quad (3.305)$$

Hieraus ergeben sich direkt die Lösungen: $17x_3 = 17 \Rightarrow x_3 = 1$, $2x_2 + 5x_3 = 9 \Rightarrow x_2 = 4/2 = 2$ und $2x_1 - 2 + 3 = 2x_1 + 1 = 1 \Rightarrow x_1 = 0$.

Man kann die ursprüngliche Matrix A und die transformierte Matrix $\tilde{A}$ einander gegenüber stellen:

$$\begin{pmatrix} 2 & -1 & 3 \\ 3 & 1 & -2 \\ 1 & 1 & 1 \end{pmatrix}, \quad \begin{pmatrix} 2 & -1 & 3 \\ 0 & 2 & 5 \\ 0 & 0 & 17 \end{pmatrix}$$

und sieht die Transformation von A in eine Dreiecksmatrix direkt. Die hier vorgeführte Transformation hat einen gewissen ad hoc-Charakter und mag insofern nicht zufriedenstellend erscheinen, und natürlich existieren kanonische Darstellungen des Algorithmus (Golub & van Loan (2013), Kapitel 3), auf die hier aber nicht eingegangen werden soll bzw. kann, da diese Darstellungen auf der einen Seite recht lang sind und auf der anderen nicht benötigt werden, wenn man nur an der Anwendung interessiert ist, was im Allgemeinen der Fall sein wird; Statistikpakete enthalten entsprechende Module. Hier soll nur festgestellt werden, dass A und $\tilde{A}$ natürlich denselben Rang haben.

Im folgenden Abschnitt wird anhand der SVD von A der Teilraum $\mathbf{L}_{n-r}$ näher bestimmt.

3.16.4 Die Cholesky-Zerlegung

Die Anwendung des Gauß-Algorithmus impliziert, dass die Koeffizientenmatrix A des Gleichungssystems $A\mathbf{x} = \mathbf{y}$ in eine Dreiecksform überführt wird. Für den Spezialfall einer symmetrischen, positiv-definiten Matrix läßt sich zeigen, dass A stets als Produkt zweier Dreiecksmatrizen darstellbar ist; dieser Sachverhalt bedeutet, dass der Aufwand für das Lösen des Gleichungssystems drastisch reduziert wird; dies ist die Cholesky-Zerlegung²⁵. In diesem Skript wird auf numerische Fragen kaum eingegangen, da die Programmpakete für multivariate Verfahren im Allgemeinen effiziente Algorithmen enthalten, um die sich der Anwender nicht

²⁵ André-Louis Cholesky (1875 – 1918), französischer Mathematiker

weiter kümmern muß. In einigen theoretischen Herleitungen wird aber auf die Cholesky-Zerlegung eingegangen, so dass in diesem Absatz kurz auf sie eingegangen werden soll.

Satz 3.29 *Es sei A eine symmetrische, positiv-definite Matrix. Dann existiert eine untere Dreiecksmatrix L derart, dass*

$$A = LL'. \quad (3.306)$$

Beweis: Statt eines allgemeinen Beweises wird der Satz anhand einer (3×3) -Matrix illustriert; das angewendete Prinzip überträgt sich sofort auf den allgemeinen $(n \times n)$ -Fall. Nach Behauptung existieren also L und L^* derart, dass

$$A = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix} = \begin{pmatrix} L_{11} & 0 & 0 \\ L_{21} & L_{22} & 0 \\ L_{31} & L_{32} & L_{33} \end{pmatrix} \begin{pmatrix} L_{11} & L_{21} & L_{31} \\ 0 & L_{22} & L_{32} \\ 0 & 0 & L_{33} \end{pmatrix}$$

Rechnet man das Produkt auf der rechten Seite aus, so erhält man

$$A = LL' = \begin{pmatrix} L_{11}^2 & L_{11}L_{21} & L_{11}L_{31} \\ L_{21}L_{11} & L_{21}^2 + L_{22}^2 & L_{21}L_{31} + L_{22}L_{32} \\ L_{31}L_{11} & L_{31}L_{21} + L_{32}L_{22} & L_{31}^2 + L_{32}^2 + L_{33}^2 \end{pmatrix}.$$

Das Element a_{ij} von A ist dann gleich dem Element in der i -ten Zeile und j -ten Spalte der Matrix auf der rechten Seite, also

$$a_{11} = L_{11}^2, \quad a_{22} = L_{21}^2 + L_{22}^2, \quad a_{33} = L_{31}^2 + L_{32}^2 + L_{33}^2,$$

und

$$a_{21} = L_{21}L_{11}, \quad a_{23} = L_{21}L_{31} + L_{22}L_{32}, \quad a_{31} = L_{31}L_{11}$$

und wegen der vorausgesetzten Symmetrie von A gilt $a_{ij} = a_{ji}$, so dass nicht mehr Elemente bestimmt werden müssen. Wegen der vorausgesetzten Positiv-Definitheit gilt $a_{11} > 0$, also folgt $L_{11} = \sqrt{a_{11}}$. Dann folgt $L_{21} = a_{21}/L_{11} = a_{21}/\sqrt{a_{11}}$ und $L_{31} = a_{31}/\sqrt{a_{11}}$, etc. Auf diese Weise lassen sich die L_{ij} aus den a_{ij} berechnen. Das Prinzip kann leicht auf den Fall allgemeiner $(n \times n)$ -Matrizen übertragen werden. $\square$

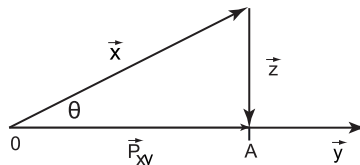
Die LDL-Zerlegung: Eine etwas verallgemeinerte Form der Cholesky-Zerlegung ist die *LDL-Zerlegung* einer symmetrischen, positiv-definiten Matrix:

$$A = LDL', \quad (3.307)$$

wobei L eine untere Dreiecksmatrix und D eine Diagonalmatrix ist. Man kann $D = D^{1/2}D^{1/2}$ schreiben, wobei $D^{1/2}$ eine Diagonalmatrix ist, deren Diagonalelemente die Wurzeln aus den Diagonalelementen von D ist, d.h. die Diagonalelemente von D müssen größer als Null sein. Setzt man $G = LD^{1/2}$, so hat man

$$A = GG'. \quad (3.308)$$

Abbildung 12: Orthogonale Projektion des Vektors $\mathbf{x}$ auf einen Vektor $\mathbf{y}$ bzw. auf eine Gerade



Es sei $A\mathbf{x} = \mathbf{y}$ ein Gleichungssystem. Setzt man $G\mathbf{z} = \mathbf{y}$, $G'\mathbf{x} = \mathbf{z}$, so hat man $GG'\mathbf{x} = \mathbf{y}$. Man löst also das Gleichungssystem, indem man zunächst $G\mathbf{z} = \mathbf{y}$ und dann $G'\mathbf{x} = \mathbf{z}$ löst; da G und G' in Dreiecksform sind, sind diese beiden Gleichungssysteme schnell und einfach zu lösen. Man findet für die Elemente g_{ik} von G

$$g_{ik} = \begin{cases} 0, & i < k \\ \sqrt{a_{ik} - \sum_{j=1}^{k-1} g_{kj}^2}, & i = k \\ \frac{1}{a_{ik}} \left(a_{ik} - \sum_{j=1}^{k-1} g_{ij}g_{kj} \right), & i > k \end{cases} \quad (3.309)$$

Eine ausführliche Darstellung von Lösungen von linearen Gleichungssystemen durch Rückführung auf Dreiecksmatrizen wird in Golub & vanLoan (2013), Kap. 3 und 4, gegeben.

3.17 Projektionen

3.17.1 Orthogonale Projektion eines Vektors auf einen anderen

Die Projektion eines Vektors auf einen anderen spielt in vielen Anwendungen eine wichtige Rolle. Um die Idee zu illustrieren, wird die Projektion eines Vektors $\mathbf{y}$ auf einen Vektor $\mathbf{x}$ betrachtet, vergl. Abbildung 12. $\mathbf{x}$ und $\mathbf{y}$ schließen den Winkel θ ein, und es ist

$$\vec{P}_{yx} = a\mathbf{y} \quad (3.310)$$

der Vektor, der sich ergibt, wenn $\mathbf{x}$ auf $\mathbf{y}$ projiziert wird; in Abbildung 12 ist $a < 1$, aber $a > 1$ ist möglich. Nun ist einerseits

$$\mathbf{z} = \vec{P}_{yx} - \mathbf{x} = a\mathbf{y} - \mathbf{x}, \quad (3.311)$$

und da $\mathbf{z}$ andererseits senkrecht auf $\mathbf{y}$ steht muß $\mathbf{z}'\vec{P}_{yx} = (\vec{P}_{yx} - \mathbf{x})'\vec{P}_{yx} = 0$ gelten, so dass $\vec{P}_{yx}'\vec{P}_{yx} = \mathbf{x}'\vec{P}_{yx}$ gelten muß, d.h. es gilt $a^2\|\mathbf{y}\|^2 = \mathbf{x}'\mathbf{y}$, so dass

$$a = \frac{\mathbf{x}'\mathbf{y}}{\|\mathbf{y}\|^2} \Rightarrow \vec{P}_{xy} = \frac{\mathbf{x}'\mathbf{y}}{\|\mathbf{y}\|^2}\mathbf{y} \quad (3.312)$$

folgt.

Man kann die Länge von $\vec{P}_{xy}$ als Koordinate der Projektion des Endpunkts von $\mathbf{x}$ auf eine Koordinatenachse $\mathbf{y}$ interpretieren. Diese Koordinate ist dann durch

$$\|\vec{P}_{xy}\| = a\|\mathbf{y}\| \quad (3.313)$$

gegeben. Im folgenden Abschnitt wird dieser Aspekt von Projektionen weiter elaboriert.

3.17.2 Projektionen auf Hauptachsen

Es sei X eine $(m \times n)$ -Matrix von Messwerten; x_{ij} sei der Messwert des i -ten Objects ("Person") für die j -te Variable ("Test"), $1 \leq i \leq m$ und $1 \leq j \leq n$. Für X gilt die Singularwertzerlegung $X = Q\Lambda^{1/2}T' = LP'$ mit $L = Q\Lambda^{1/2}$. Wegen der Orthonormalität von T folgt $XT = L$. Das Element ℓ_{ik} von L ist die Koordinate des i -ten Falls auf der k -ten latenten Dimension und ergibt sich als Skalarprodukt der i -ten Zeile ξ'_i von X und der k -ten Spalte $\mathbf{t}_k$ von T , d.h. es ist

$$\ell_{ik} = \xi'_i \mathbf{t}_k. \quad (3.314)$$

Hier wird der Zeilenvektor als *Spaltenvektor* aufgefasst, um die Schreibweise des Skalarprodukts beizubehalten. $\mathbf{t}_k$ definiert die Orientierung der k -ten Hauptachse eines Ellipsoids, und ℓ_{ik} ist die Länge des Vektors $\vec{P}_{ik}$, der sich als Projektion von ξ_i auf die k -te Hauptachse des Ellipsoids ergibt, das die Punktekonfiguration der Fälle repräsentiert. $\vec{P}_{ik}$ entspricht dem Vektor $\vec{P}_{xy}$ in Abbildung 12. Dem vorangegangenen Abschnitt zufolge kann $\vec{P}_{ik} = a_{ik}\mathbf{t}_k$, $a_{ik} \in \mathbb{R}$, geschrieben werden, wobei der Faktor a_{ik} indiziert wurde um anzuzeigen, dass er für ξ_i und $\mathbf{t}_k$ charakteristisch ist. Nach (3.312) gilt nun

$$a_{ik} = \frac{\xi'_i \mathbf{t}_k}{\|\mathbf{t}_k\|^2} = \xi'_i \mathbf{t}_k, \quad (3.315)$$

da ja $\|\mathbf{t}_k\| = 1$, und nach (3.313) hat man dann

$$\|\vec{P}_{ik}\| = a_{ik}\|\mathbf{t}_k\| = \xi'_i \mathbf{t}_k, \quad (3.316)$$

d.h. wegen (3.314) gilt

$$\ell_{ik} = \|\vec{P}_{ik}\|, \quad (3.317)$$

so dass die Koordinate ℓ_{ik} durch die Länge der Projektion des Vektors ξ_i auf die k -te Hauptachse des Ellipsoids gegeben ist.

3.17.3 Projektionen auf k -dimensionale Teilräume

Die Projektion von n -dimensionalen Vektoren auf einen 1-dimensionalen Teilraum ist ein Spezialfall. So versucht man in der Diskriminanzanalyse, eine mehrdimensionale Punktekonfiguration auf einen möglichst niedrigdimensionalen Teilraum

zu projizieren, in dem bestimmte Klassen von Punkten optimal separiert repräsentiert werden, – ein 1-dimensionaler Teilraum ist u. U. nicht hinreichend, um die Klassifikationsaufgabe hinreichend gut zu lösen.

So wird jetzt nicht ein einzelner Vektor $\mathbf{v}$ betrachtet, der einen 1-dimensionalen Teilraum definiert, sondern ein k -dimensionaler Teilraum V des $\mathbb{R}^m$, $k < n$, mit $W = V^\perp$ als seinem orthogonalen Komplement. $V^\perp$ ist ebenfalls ein Teilraum.

Es sei $\mathbf{x} \in \mathbb{R}^m$, also ein m -dimensionaler Vektor. Nach Satz 2.15, Seite 56 kann dann

$$\mathbf{x} = \mathbf{v} + \mathbf{w}, \quad \mathbf{v} \in V, \quad \mathbf{w} \in W \quad (3.318)$$

mit $\mathbf{v} \in \mathbb{R}^m$, $\mathbf{w} \in \mathbb{R}^m$ geschrieben werden. $\mathbf{v} = P_V(\mathbf{x})$ ist die Projektion von $\mathbf{x}$ auf V , und $\mathbf{w} = P_W(\mathbf{x})$ ist die Projektion von $\mathbf{x}$ auf das orthogonale Komplement W , so dass man auch

$$\mathbf{x} = P_V(\mathbf{x}) + P_W(\mathbf{x}) \quad (3.319)$$

schreiben kann. Für V kann eine orthogonale Basis $(\mathbf{a}_1, \dots, \mathbf{a}_k)$ angenommen werden; die $\mathbf{a}_j$ können als Spaltenvektoren einer Matrix A zusammengefasst werden:

$$A = [\mathbf{a}_1 | \dots | \mathbf{a}_k]. \quad (3.320)$$

Da $\mathbf{v} \in V$, existiert dann ein k -dimensionaler Vektor $\mathbf{y} \in \mathbb{K}$ (d.h. hier $\mathbf{y} \in \mathbb{R}^k$) derart, dass

$$\mathbf{v} = y_1 \mathbf{a}_1 + \dots + y_k \mathbf{a}_k = A\mathbf{y}. \quad (3.321)$$

Wegen $\mathbf{x} = \mathbf{v} + \mathbf{w}$ folgt $\mathbf{w} = \mathbf{x} - \mathbf{v} \in W = V^\perp$. Es ist

$$\mathcal{L}(A) = \mathcal{L}(\mathbf{a}_1, \dots, \mathbf{a}_k) = V.$$

Nun ist aber $\mathbf{a}_k' \mathbf{w} = 0$ für alle $\mathbf{a}_k$, denn $\mathbf{w}$ ist ja aus dem orthogonalen Komplement von V , so dass man allgemein

$$A' \mathbf{w} = A'(\mathbf{x} - \mathbf{v}) = \vec{0}$$

schreiben kann, denn die erste Zeile von A' enthält ja gerade $\mathbf{a}_1'$, die zweite $\mathbf{a}_2'$, etc. Dies wiederum bedeutet, dass

$$\text{kern}(A') = V^\perp = W,$$

d.h. der Kern (null space) von A' ist gerade W . Ausmultipliziert ergibt sich wegen (3.321)

$$A' \mathbf{x} - A' \mathbf{v} = A' \mathbf{x} - A' A \mathbf{y} = \vec{0} \Rightarrow A' \mathbf{x} = A' A \mathbf{y}.$$

Da die Spaltenvektoren von A orthogonal zueinander sind, hat A und damit $A' A$ den Rang k , so dass die zu $A' A$ inverse Matrix $(A' A)^{-1}$ existiert, woraus

$$(A' A)^{-1} A' \mathbf{x} = \mathbf{y} \quad (3.322)$$

folgt. Nun ist oben für den unbekannten Projektionsvektor $\mathbf{v}$ die Gleichung $A\mathbf{y} = \mathbf{v}$ aufgestellt worden; multipliziert man also (3.322) von links mit A , so erhält man

$$A(A'A)^{-1}A'\mathbf{x} = A\mathbf{y} = \mathbf{v}. \quad (3.323)$$

Man benötigt also keine explizite Lösung für den Vektor $\mathbf{y}$, sondern nur eine für die Matrix A , so dass sich $P = A(A'A)^{-1}A'$ berechnen lässt. Die Matrix P hat eine Reihe von Eigenschaften, die nachgewiesen werden sollen.

Satz 3.30 *Es gelten die Aussagen*

Dann gilt

(1) *die Projektion ist eine lineare Transformation,*

(2) *P ist symmetrisch,*

(3) *P ist idempotent, d.h. $P^2 = P$.*

(4) *Für $\mathbf{v} \in \mathcal{L}(A)$ gilt $P\mathbf{v} = \mathbf{v}$.*

Beweis: (1) folgt sofort aus der Gleichung $P\mathbf{x} = \mathbf{v}$, die ja eine lineare Transformation von $\mathbf{x}$ in $\mathbf{v}$ repräsentiert. Für (2) findet man sofort

$$P' = (A(A'A)^{-1}A')' = A(A'A)^{-1}A' = P,$$

also ist P symmetrisch, und (3) gilt wegen

$$\begin{aligned} P^2 &= (A(A'A)^{-1}A')(A(A'A)^{-1}A') \\ &= A(A'A)^{-1}A'A(A'A)^{-1}A' = A(A'A)^{-1}A' = P \end{aligned}$$

(4): $\mathbf{v} \in \mathcal{L}(A)$ impliziert die Existenz eines Vektors $\mathbf{y}$ mit $A\mathbf{y} = \mathbf{v}$, d.h. $\mathbf{v}$ ist eine Linearkombination der Spalten von A . Dann folgt

$$P\mathbf{v} = A(A'A)^{-1}A'\mathbf{v} = A(A'A)^{-1}A'A\mathbf{y} = A\mathbf{y} = \mathbf{v}.$$

□

Bei der in (3.320) definierten Matrix ist vorausgesetzt worden, dass ihre Spalten orthogonal sind. Diese Annahme impliziert, dass $(A'A)^{-1}$ stets existiert, da A und die $(k \times k)$ -Matrix $A'A$ dann den vollen Rang k haben. Ist die Basis insbesondere orthonormal, so ist $A'A = I$ die Einheitsmatrix und (3.323) vereinfacht sich zu

$$AA'\mathbf{x} = \mathbf{v}. \quad (3.324)$$

Es zeigt sich aber, dass (3.323) verallgemeinert werden kann für den Fall, dass die Inverse $(A'A)^{-1}$ nicht existiert; in diesem Fall kann die generalisierte Inverse $(A'A)^{-}$ eingesetzt werden, so dass

$$A(A'A)^{-}A'\mathbf{x} = A\mathbf{y} = \mathbf{v}. \quad (3.325)$$

gilt. Die in Satz 3.30 aufgeführten Eigenschaften übertragen sich auf die Matrix $A(A'A)^{-}A'$, wie man sofort nachprüft: Da $A'A$ symmetrisch ist, ist auch $(A'A)^{-}$ symmetrisch, so dass

$$(A(A'A)^{-}A')' = A(A'A)^{-}A'.$$

Weiter ist $A(A'A)^-A'$ idempotent, denn

$$(A(A'A)^-A')(A(A'A)^-A') = A(A'A)^-A',$$

und es gilt für $\mathbf{v} \in \mathcal{L}(A)$, d.h. $\mathbf{v} = A\mathbf{y}$

$$A(A'A)^-A'\mathbf{v} = A(A'A)^-A'A\mathbf{y} = A\mathbf{y} = \mathbf{v}.$$

Man hat dann die folgende

Definition 3.23 *Es sei A eine $(n \times k)$ -Matrix. Dann heißt*

$$H = A(A'A)^-A' \quad (3.326)$$

einschließlich des Spezialfalls $(A'A)^- = (A'A)^{-1}$ Projektionsmatrix oder Projektionsoperator (projection matrix, hat-matrix).

Satz 3.31 *Es sei X eine $(m \times n)$ -Matrix. Dann ist $X(X'X)^-X'$ ein Projektionsoperator.*

Beweis: Es sei $\mathbf{u} \in \mathcal{L}(X)$; dann gilt $\mathbf{u} = \mathbf{v} + \mathbf{w}$ mit $\mathbf{v} \in \mathcal{L}(X)$ und $\mathbf{w} \in \mathcal{L}(X)^\perp$ (Satz 2.15, Gleichung (2.98), Seite 56). Dann ist $\mathbf{v} = X\mathbf{b}$ und

$$X(X'X)^-X'\mathbf{v} = X(X'X)^-X'X\mathbf{b} = X\mathbf{b} = \mathbf{v}, \quad X(X'X)^-X'\mathbf{w} = 0$$

weil $X'\mathbf{w} = 0$, da ja $\mathbf{w} \perp X$. Es sei weiter $\mathbf{u} = X\mathbf{y}$ für einen geeignet gewählten Vektor $\mathbf{y}$. Dann ist

$$\begin{aligned} X(X'X)^-X'\mathbf{u} &= X(X'X)^-X'X\mathbf{y} = X\mathbf{y} = X(X'X)^-X'(\mathbf{v} + \mathbf{w}) \\ &= X(X'X)^-X'\mathbf{v} = \mathbf{v}. \end{aligned}$$

□

3.17.4 Projektion eines Datenvektors auf einen Teilraum

Es sei X eine $(m \times n)$ -Matrix von Messungen von n Variablen an m Fällen. Die SVD von X ist durch $X = Q\Lambda^{1/2}P'$ gegeben. Speziell für den j -ten Spaltenvektor $\mathbf{x}_j$ gilt dann $\mathbf{x}_j = Q\Lambda^{1/2}\mathbf{p}'_j$, wobei $\mathbf{p}'_j$ der j -te Spaltenvektor von P' ist, d.h. $\mathbf{p}_j$ ist der j -te Zeilenvektor von P . Q ist die $(m \times n)$ -Matrix der orthonormalen Eigenvektoren von XX' , P ist die $(n \times n)$ -Matrix der orthonormalen Eigenvektoren von $X'X$, und Λ ist die $(n \times n)$ -Diagonalmatrix der Eigenwerte von XX' bzw. $X'X$.

Es sei Q_k die Teilmatrix der ersten $k < n$ Spaltenvektoren von Q . Diese Vektoren spannen einen k -dimensionalen Teilraum des $\mathbb{R}^m$ auf. Gesucht ist die Projektion von $\mathbf{x}_j$ auf $\mathcal{L}(Q_k)$. Es ist

$$\mathbf{x}_j = \mathbf{v}_j + \mathbf{w}_j \quad (3.327)$$

mit $\mathbf{v}_j \in \mathcal{L}(Q_k)$ und $\mathbf{w}_j \in \mathcal{L}(Q_k)^\perp$; $\mathbf{v}_j$ ist die Projektion von $\mathbf{x}_j$ auf $\mathcal{L}(Q_k)$. Es existiert dann ein Vektor $\mathbf{y}$ derart, dass $\mathbf{v}_j = Q_k \mathbf{y}$. Dann ist

$$Q'_k \mathbf{v}_j = Q'_k \mathbf{v}_j + Q'_k \mathbf{w}_j = Q'_k \mathbf{v}_j,$$

denn $Q_k \mathbf{w}_j = \vec{0}$, da $\mathbf{w}_j \in \mathcal{L}(Q_k)^\perp$. Also hat man

$$Q_k \mathbf{x}_j - Q'_k \mathbf{v}_j = Q_k \mathbf{x}_j - Q'_k Q_k \mathbf{y} = \vec{0} \Rightarrow Q_k \mathbf{x}_j = Q'_k Q_k \mathbf{y},$$

so dass

$$(Q'_k Q_k)^{-1} Q'_k \mathbf{x}_j = \mathbf{y}.$$

Dann folgt durch Multiplikation von links mit Q_k

$$Q_k (Q'_k Q_k)^{-1} Q'_k \mathbf{x}_j = Q_k \mathbf{y} = \mathbf{v}_j, \quad (3.328)$$

wobei wegen der Orthonormalität der Spalten von Q_k $(Q'_k Q_k)^{-1} = I^{-1} = I$ folgt, so dass man

$$Q_k Q'_k \mathbf{x}_j = \mathbf{v}_j \quad (3.329)$$

erhält. Nun war oben aber $\mathbf{x}_j = Q \Lambda^{1/2} \mathbf{p}'_j$ festgestellt worden. Setzt man diesen Ausdruck für $\mathbf{x}_j$ ein, so erhält man

$$Q_k Q'_k Q \Lambda^{1/2} \mathbf{p}'_j = \mathbf{v}_j.$$

Wegen der Orthonormalität der Spalten von Q findet man

$$Q'_k Q = \begin{pmatrix} 1 & 0 & \cdots & 0 & 0 & \cdots & 0 \\ 0 & 1 & \cdots & 0 & 0 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots & \cdots & \vdots \\ 0 & \cdots & 0 & 1 & 0 & \cdots & 0 \end{pmatrix}.$$

$Q'_k Q$ ist eine $(k \times n)$ -Matrix. Die ersten k Zeilen und Spalten von $Q'_k Q$ enthalten eine $(k \times k)$ Einheitsmatrix und die letzten $n - k$ Spalten enthalten nur Nullen. $Q_k Q'_k Q \Lambda^{1/2}$ hat dann mit $\sigma_j = \sqrt{\lambda_j}$ die Form

$$Q_k Q'_k Q \Lambda^{1/2} = \begin{pmatrix} \sigma_1 q_{11} & \cdots & \sigma_k q_{1k} & 0 & \cdots & 0 \\ \sigma_1 q_{21} & \cdots & \sigma_k q_{2k} & 0 & \cdots & 0 \\ \vdots & \vdots & \vdots & \vdots & \ddots & \vdots \\ \sigma_1 q_{m1} & \cdots & \sigma_k q_{mk} & 0 & \cdots & 0 \end{pmatrix},$$

so dass man für die Projektion von $\mathbf{x}_j$ auf $\mathcal{L}(Q_k)$ die Beziehung

$$H \mathbf{x}_j = \mathbf{v}_j = Q_k Q'_k Q \Lambda^{1/2} \mathbf{p}'_j = \begin{pmatrix} \sum_{i=1}^k \sigma_i q_{1i} p_{ij} \\ \sum_{i=1}^k \sigma_i q_{2i} p_{ij} \\ \vdots \\ \sum_{i=1}^k \sigma_i q_{mi} p_{ij} \end{pmatrix} = \sum_{i=1}^k \sigma_i p_{ij} \mathbf{q}_i \quad (3.330)$$

erhält. Die Projektion $\mathbf{v}_j$ von $\mathbf{x}_j$ auf den durch Q_k definierten Teilraum des $\mathbb{R}^m$ ist also eine Linearkombination der k Spaltenvektoren von Q_k mit den für den j -ten Vektor spezifischen "Gewichten" $\sigma_i p_{ij}$.

Die Beziehung (3.328), also

$$Q_k(Q_k'Q_k)^{-1}Q_k'\mathbf{x}_j = Q_kQ_k'\mathbf{x}_j = \mathbf{v}_j$$

legt nahe, den Projektionsvektor $\mathbf{v}_j$ als eine "Vorhersage" von $\mathbf{x}_j$ im Sinne der Methode der Kleinsten Quadrate, also im Sinne der linearen Regression zu interpretieren, wie in Abschnitt ??, Gleichung (??) (Seite ??) gezeigt wird. Wenn man für die dortige Gleichung X der Prädiktoren die Matrix Q_k einsetzt, die hier ja die Rolle der Prädiktoren spielt. Man kann dann mit $\hat{\mathbf{x}}_j = \mathbf{v}_j$

$$\mathbf{x}_j = \hat{\mathbf{x}}_j + \hat{\mathbf{e}}_j \quad (3.331)$$

schreiben. $\hat{\mathbf{e}}_j$ repräsentiert dann alle Einflüsse, die nicht durch die ersten k Spaltenvektoren von Q und P in der SVD $X = Q\Lambda^{1/2}P'$ erklärt werden. Setzt man

$$\hat{X}_k = [\hat{\mathbf{x}}_1 | \dots | \hat{\mathbf{x}}_k], \quad (3.332)$$

ist $\hat{X}$ also eine $(m \times k)$ -Matrix, deren Spaltenvektoren die $\hat{\mathbf{x}}_j$, $j = 1, \dots, k$ sind, so kann man allgemein

$$X = \hat{X}_k + \hat{E}_k \quad (3.333)$$

schreiben, wobei die Spaltenvektoren von $\hat{E}_k$ die Fehlervektoren $\hat{\mathbf{e}}_j$ aus (3.331) sind. Dann ist $X - \hat{X}_k = \hat{E}_k$ und

$$(X - \hat{X}_k)'(X - \hat{X}_k) = \|X - \hat{X}_k\|^2 = \hat{E}_k'\hat{E}_k, \quad (3.334)$$

und wegen Satz 3.24, Gleichung (3.244) (Seite 133) gilt

$$\|X - \hat{X}_k\|^2 = \lambda_{k+1}, \quad (3.335)$$

und nach Satz 3.25 (Seite 134) ist dies die beste Approximation von X im Sinne des kleinstmöglichen Wertes von $\|X - \hat{X}_k\|$ auf der Basis von k latenten Vektoren.

3.18 Schlecht konditionierte Matrizen und Regularisierung*

Die Methode der Kleinsten Quadrate (KQ-Methode) führt bei der linearen Regression auf das Gleichungssystem

$$(X'X)\hat{\mathbf{b}} = X'\mathbf{y}$$

mit dem unbekannten Vektor $\hat{\mathbf{b}}$. Die Messungen der Prädiktoren in der Matrix X und der abhängigen Variablen im Vektor $\mathbf{y}$ sind üblicherweise nicht fehlerfrei, und diese Ungenauigkeiten zusammen mit eventuellen Kollinearitäten bzw. Multikollinearitäten zwischen den Prädiktoren wirken sich auf die Genauigkeit, mit der $\hat{\mathbf{b}}$

den Vektor $\mathbf{b}$ der Regressionskoeffizienten schätzt, aus. Das Ausmaß, in dem sich Ungenauigkeiten auf die Lösung auswirken, wird als *Sensitivität* bezeichnet. Hat man also allgemein das Gleichungssystem $A\mathbf{x} = \mathbf{y}$ und hat die $(n \times n)$ -Matrix A vollen Rang (ist nicht singulär), so erhält man die Lösung durch $\mathbf{x} = A^{-1}\mathbf{y}$. Für A kann man die SVD anschreiben:

$$A = Q\Lambda^{1/2}P', \quad (3.336)$$

wobei Q , Λ und P $(n \times n)$ -Matrizen sind und wegen des vollen Rangs von A alle Eigenwerte in Λ von Null verschieden sind (man bemerke, dass im allgemeinen Fall A nicht notwendig symmetrisch ist). Für die Inverse A^{-1} erhält man den Ausdruck

$$A^{-1} = (Q\Lambda^{1/2}P')^{-1} = P\Lambda^{-1/2}Q' = \sum_{j=1}^n \frac{\mathbf{q}_j\mathbf{p}_j'}{\lambda_j}, \quad (3.337)$$

denn Q und P sind ja orthonormal, so dass $Q^{-1} = Q'$, $P^{-1} = P'$. Einige der Implikationen des Falles, dass einige λ_j klein werden, sind in Abschnitt ?? bereits angesprochen worden. Kleine Veränderungen ("Störungen") in A oder $\mathbf{y}$ können große Veränderungen in $\mathbf{x}$ nach sich ziehen (s. a. Satz 3.25, Seite 134). Insbesondere werden die Elemente von A^{-1} groß, wenn λ_j -Werte klein werden. In Abschnitt 3.10 ist der Begriff der Matrixnorm eingeführt worden. Er erweist sich als nützlich, wenn man die Sensitivität, also das Ausmaß, in dem etwa die Eigenwerte einer Matrix sich auf die Inverse einer Matrix auswirken, in einer Maßzahl ausdrücken will:

Definition 3.24 *Es sei A eine Matrix mit der Norm $\|A\|$. Dann heißt der Quotient*

$$\kappa(A) = \frac{\|A\|}{\|A^{-1}\|} \quad (3.338)$$

die Konditionierungszahl der Matrix A . Ist A singulär (d.h. hat A nicht den vollen Rang, so dass A^{-1} nicht existiert, wird $\kappa(A) = \infty$ gesetzt.

Die Matrixnorm sei etwa durch die Frobenius-Norm (3.236) gegeben. Der Wert von $\|A\|$ ist dann durch die Summe der $|a_{ij}|^2$ gegeben. Ist A schlecht konditioniert, so sind die Elemente der Inversen A^{-1} groß im Vergleich zu den a_{ij} und die Norm von A^{-1} ist groß im Vergleich zur Norm $\|A\|$. Für symmetrische Matrizen ist $\kappa(A)$ durch

$$\kappa(A) = \frac{\lambda_{\max}}{\lambda_{\min}} \quad (3.339)$$

gegeben, $\lambda_{\max}$ der größte und $\lambda_{\min}$ der kleinste Eigenwert von A . Je kleiner also der kleinste Eigenwert der Matrix, desto größer die Konditionierungszahl. Der folgende Sachverhalt erläutert die Bedeutung der Konditionierungszahl.

Insbesondere Gleichungssysteme mit einer großen Anzahl von Unbekannten werden iterativ gelöst (auf die Details der numerischen Verfahren wird hier nicht

eingegangen). Dabei spielen Rundungsfehler eine Rolle, die eine Abweichung von der wahren Lösung bewirken. Die Elemente der Matrix A des Gleichungssystems sind im Allgemeinen reelle Zahlen, die als Vereinigung der rationalen und der irrationalen Zahlen definiert sind. Rationale Zahlen sind als Quotient (lat. *ratio*) p/q darstellbar, wobei p, q natürliche Zahlen sind, also $p, q \in \mathbb{N}$. Bei rationalen Zahlen ist die Anzahl der Zahlen nach dem Dezimalpunkt (Dezimalkomma) entweder endlich oder periodisch. Irrationalzahlen sind nicht unvernünftig, sondern einfach *nicht* als Quotient – daher irrational, also nicht als ratio – von natürlichen Zahlen darstellbar, und sie haben unendlich viele Stellen nach dem Dezimalpunkt und es existiert keine Periode. Es gibt wesentlich mehr irrationale Zahlen als rationale Zahlen (auf eine genaue Darstellung der Mächtigkeitsverhältnisse muß hier verzichtet werden), – man kann sagen, dass man es im Allgemeinen mit Irrationalzahlen zu tun hat²⁶. Bei jeder tatsächlich durchgeführten Berechnung muß man sich sich notwendig auf eine endliche Anzahl von Nachkommastellen beschränken. Dadurch entsteht ein *Anfangsabbrechfehler*. Die Abweichung der Rechnung vom wahren Wert ist von der Größenordnung von $\kappa(A)$ Einheiten der letzten Dezimalstelle.

3.19 Kroneckerprodukte

Es sei A eine $m \times n$ -Matrix und B eine $p \times q$ -Matrix. Dann ist das *Kroneckerprodukt* von A und B durch die Matrix

$$C = A \otimes B = \begin{pmatrix} a_{11}B & a_{12}B & \cdots & a_{1n}B \\ a_{21}B & a_{22}B & \cdots & a_{2n}B \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1}B & a_{m2}B & \cdots & a_{mn}B \end{pmatrix} \quad (3.340)$$

definiert; das Zeichen $\otimes$ signalisiert die Bildung des Kroneckerprodukts.

X sei eine $M \times n$ -Matrix. Man kann aus X Vektoren bilden: einmal, indem man die Zeilenvektoren aneinander reiht. Es entsteht dann ein mn -dimensionaler Zeilenvektor $\mathbf{x}^z$. Man kann andererseits alle Spaltenvektoren von X untereinander anschreiben. Es entsteht ein nm -dimensionaler Spaltenvektor $\mathbf{x}^v$. Man spricht dann von der *Vektorisierung* der Matrix X . Das Kroneckerprodukt $A \otimes B$ hat

²⁶Es sei X eine auf einem reellen Intervall definierte zufällige Veränderliche. Es läßt sich zeigen, dass die Wahrscheinlichkeit, dass X einen rationalen Wert annimmt gleich Null ist!

dann die folgenden Eigenschaften:

$$\lambda(A \otimes B) = (\lambda A) \otimes B = A \otimes (\lambda B), \lambda \in \mathbb{R} \quad (3.341)$$

$$A \otimes (B \otimes C) = (A \otimes B) \otimes C = A \otimes B \otimes C \quad (3.342)$$

$$(A \otimes B)' = A' \otimes B' \quad (3.343)$$

$$(A \otimes B)(F \otimes G) = (AF) \otimes (BG) \quad (3.344)$$

$$(A \otimes B)^{-1} = A^{-1} \otimes B^{-1} \quad (3.345)$$

$$(A + B) \otimes C = A \otimes C + B \otimes C \quad (3.346)$$

$$A \otimes (B + C) = A \otimes B + A \otimes C \quad (3.347)$$

$$(AXV)^v = (B' \otimes A)\mathbf{x}^v \quad (3.348)$$

4 Abbildungen und Funktionen

4.1 Allgemeine Definition von Abbildungen

Es werde zunächst der allgemeine Begriff der Abbildung erklärt. Eine Abbildung f einer Menge X in eine Menge Y ordnet jedem Element genau einem Element aus Y zu:

$$f : X \rightarrow Y, \quad x \mapsto y = f(x). \quad (4.1)$$

$f : X \rightarrow Y$ ist die Beziehung zwischen den Mengen X und Y , die durch f erklärt wird, und $x \mapsto y = f(x)$ gibt f als Zuordnung des Elements $x \in X$ zum Element $y = f(x) \in Y$ an. Wichtig ist hierbei, dass einem Element $x \in X$ nur ein Element $y \in Y$ zugeordnet wird. Dies schließt nicht aus, dass verschiedenen Elementen $x \in X$ der gleiche Wert $y \in Y$ zugeordnet werden kann. In diesem Fall kann von einem Element $y \in Y$ nicht eindeutig auf das Element $x \in X$ mit $f(x) = y$ zurückgeschlossen werden.

Mit der Schreibweise $f(X)$ ist nicht ein einzelnes Element gemeint, sondern die Menge der Werte, die man erhält, wenn man f für alle Werte aus X bestimmt, also

$$f(X) = \{f(x), x \in X\}. \quad (4.2)$$

Offenbar gilt $f(X) \subseteq Y$. $f(X)$ heißt das *Bild von X in Y* , und X ist das *Urbild von $f(X)$* .

Definition 4.1 *Es sei $f : X \rightarrow Y$. Dann ist f*

1. *injektiv, wenn aus $x, x' \in X$ und $f(x) = f(x')$ folgt, dass $x = x'$ (und damit $f(x) \neq f(x') \Rightarrow x \neq x'$). Es kann $f(X) \subset Y$ gelten, d.h. $f(X)$ kann eine echte Teilmenge von Y sein.*
2. *surjektiv, wenn $f(X) = f(Y)$, d.h. zu jedem $y \in Y$ existiert ein $x \in X$ derart, dass $y = f(x)$. Es gilt $f(X) = Y$.*

3. bijektiv, wenn f sowohl injektiv als auch surjektiv ist. Es gilt $f(X) = Y$.

Beispiele: $f : \mathbb{R} \rightarrow \mathbb{R}$, $x \mapsto ax + b$ für $a, b \in \mathbb{R}$ fest gewählte Konstante. f ist sicher injektiv, denn $f(x) = f(x')$ impliziert $ax + b = ax' + b$ und damit $x = x'$, wie man leicht nachrechnet. f ist auch surjektiv, denn für $y = ax + b$ existiert genau ein $x = (y - b)/a$ derart, dass $y = f(x)$. Da f sowohl injektiv wie surjektiv ist, ist f auch bijektiv.

Nun sei $f^2(x) = ax$ für ein $a \in \mathbb{R}$. f ist *keine* Funktion, denn $x \in \mathbb{R}$ soll nur *ein* y -Wert zugeordnet werden. Tatsächlich wird aber x zwei Werten, nämlich $-y$ und y zugeordnet, denn $(-y)^2 = y^2$ erfüllen die Bedingung $f^2 = ax$. Zur Spezifikation der Funktion muß also angegeben werden, welchem y -Wert x zugeordnet werden soll.

Nun sei $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $y = f(x) = x^2$. Es sei $f(x) = f(x')$. Dann folgt nicht $x = x'$, denn für $x \neq x'$ gilt ebenfalls $f(x) = f(x')$, nämlich wenn $x' = -x$. Also ist f nicht injektiv. f ist auch nicht surjektiv. Denn es sei $x^2 = -1$, so dass $x = \sqrt{-1}$. Aber $\sqrt{-1}$ ist keine reelle Zahl. Dann ist f natürlich auch nicht bijektiv. Nun sei $\mathbb{R}_+$ die Menge der positiven reellen Zahlen, $f : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ und $f(x) = x^2$. Es gibt stets eine positive Zahl x , also $x \in \mathbb{R}_+$, derart dass $x^2 \in \mathbb{R}_+$. Jetzt ist f surjektiv und injektiv, also auch bijektiv.

Es sei $\mathbb{N} = \{1, 2, 3, \dots\}$ die Menge der natürlichen Zahlen, und $\mathbb{N} \cup \{0\}$ sei die Menge der natürlichen Zahlen, zu der die Null hinzugenommen worden sei. Dann ist $F : \mathbb{N} \rightarrow \mathbb{N} \cup \{0\}$ mit $f(n) = n - 1$, $n \in \mathbb{N}$, eine bijektive Funktion. Um das zu sehen muß gezeigt werden, dass f sowohl injektiv wie surjektiv ist. Um zu sehen, dass f injektiv ist, muß man zeigen, dass aus $f(m) = f(n)$ folgt, dass $m = n \in \mathbb{N}$. Man sieht sofort

$$n - 1 = f(n) = f(m) = m - 1 \Rightarrow n - 1 = m - 1 \Rightarrow m = n,$$

also ist f injektiv. Um zu sehen, dass f auch surjektiv ist, muß man zeigen, dass für $n - 1 \in \mathbb{N} \cup \{0\}$ mindestens ein $m \in \mathbb{N}$ existiert derart, dass $f(m) = n - 1$. Sicherlich folgt aus $m \in \mathbb{N} \cup \{0\}$, dass $m + 1 \in \mathbb{N}$ (für $m = 0 \notin \mathbb{N}$ folgt $0 + 1 = 1 \in \mathbb{N}$, und dies gilt sicher für alle $0 < m \in \mathbb{N}$). Also folgt für $n = m + 1$ dann

$$f(n) = f(m + 1) = (m + 1) - 1 = m = n - 1 \in \mathbb{N} \cup \{0\}.$$

Mithin ist f surjektiv, und damit ist f dann auch bijektiv.

Abbildungen lassen sich verknüpfen:

Definition 4.2 Ist f eine Abbildung der Menge X in eine Menge Y und g eine Abbildung der Menge Y in die Menge Z , so heißt die Abbildung

$$f \circ g : X \rightarrow Z, \quad x \mapsto g(f(x)) = (g \circ f)(x) \quad (4.3)$$

die Verknüpfung, oder Komposition oder auch Hintereinanderschaltung der Abbildungen f und g .

Für eine Verknüpfung $f \circ g$ wird auch

$$X \xrightarrow{f} Y \xrightarrow{g} Z \quad (4.4)$$

geschrieben. Verknüpfungen sind assoziativ, d.h. es gilt

$$(h \circ g) \circ f = h \circ (g \circ f), \quad (4.5)$$

aber im allgemeinen nicht kommutativ, d.h. im Allgemeinen hat man

$$f \circ g \neq g \circ f. \quad (4.6)$$

Der Fall $f \circ g = g \circ f$ ist nicht ausgeschlossen, ist aber ein Spezialfall.

Eine spezielle Abbildung ist

$$\text{id}_x : X \rightarrow X, \quad x \mapsto x. \quad (4.7)$$

id_x weist dem Element $x \in X$ sich selbst zu, also das identische Element aus X , zu. Man hat dann den

Satz 4.1 *Es sei $f : X \rightarrow Y$ eine Abbildung, wobei X und Y als nicht leer vorausgesetzt werden. Dann gilt*

1. *f ist genau dann injektiv, wenn eine Abbildung $g : Y \rightarrow X$ existiert derart, dass $g \circ f = \text{id}_x$*
2. *f ist genau dann surjektiv, wenn es eine Abbildung $g : Y \rightarrow X$ gibt derart, dass $f \circ g = \text{id}_x$.*
3. *f ist genau dann bijektiv, wenn es eine Abbildung $g : Y \rightarrow X$ gibt, so dass $f \circ g = \text{id}_x$ und $g \circ f = \text{id}_y$ gilt.*

Anmerkung: Man beachte, dass die Reihenfolge der Verknüpfung von f und g für injektive und surjektive Abbildungen verschieden sind.

Beweis:(1) f sei injektiv. Für $y = f(x)$ existiert dann genau ein x mit $f(x) = y$. Es sei g eine Funktion derart, dass $g(y) = x$, und für $x_0 \in X$ beliebig $g(y) = x_0$ für alle $y \in Y - f(X)$. Dann folgt $g : Y \rightarrow X$ und $g \circ f = \text{id}_x$. Umgekehrt sei $g : Y \rightarrow X$ und $g \circ f = \text{id}_x$ gegeben. Ist $f(x) = f(x')$ für $x, x' \in X$. Dann ist $x = \text{id}_x(x) = g(f(x')) = \text{id}_{x'} = x'$, so dass f injektiv sein muß.

(2) f sei surjektiv. Für jedes $y \in Y$ wird ein $x \in X$ mit $f(x) = y$ gewählt, und es sei $g(y) := x$. Dann folgt $f \circ g = \text{id}_x$. Es sei umgekehrt $g : Y \rightarrow X$ und $y \in Y$. Dann ist $y = f(g(y))$, also ist f surjektiv.

(3) f sei bijektiv. Dann sei $g := f^{-1}$; g erfüllt dann die Bedingung. Umgekehrt sei $g : Y \rightarrow X$ mit $g \circ f = \text{id}_x$ und $f \circ g = \text{id}_y$ gegeben. Dann ist f bijektiv und $g = f^{-1}$. $\square$

Der Begriff der Abbildung ist hier sehr allgemein eingeführt worden; die Mengen X und Y , für die $f : X \rightarrow Y$ definiert sein soll, wurden nicht weiter spezifiziert. Dieser Ansatz läßt die Möglichkeit offen, Produkte von Mengen zu betrachten: so seien $X_1, \dots, X_n$ Mengen, und es sei $X = X_1 \times \dots \times X_n$, also das Cartesische Produkt der $X_1, \dots, X_n$. Expliziter formuliert ist X durch

$$X = \{(x_1, \dots, x_n) | x_i \in X_i\} \quad (4.8)$$

definiert, also als die Menge der n -Tupel, die sich ergeben, wenn man jeweils ein Element aus jeder Menge X_i auswählt und damit einen Vektor $(x_1, \dots, x_n)$ erhält. Ist $X_1 = \dots = X_n = X$, so kann man

$$X_1 \times \dots \times X_n = X \times X \times \dots \times X = X^n \quad (4.9)$$

schreiben. Für $X = \mathbb{R}$ die Menge der reellen Zahlen erhält man insbesondere $\mathbb{R}^n$ als Menge der n -dimensionalen Vektoren. Man kann dann Funktionen der Art $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ betrachten, worauf im folgenden Abschnitt näher eingegangen wird.

4.2 Lineare Abbildungen

Funktionen sind als Abbildungen einer Menge X in eine Menge Y definiert worden. Vektorräume sind Mengen, so dass man auch Funktionen als Abbildungen eines Vektorraums in sich selbst oder in einen anderen definieren kann. Hier werden insbesondere lineare Abbildungen betrachtet:

Definition 4.3 *Es seien V und W Vektorräume über einem Körper $\mathbb{K}$ und $f : V \rightarrow W$ sei eine Abbildung von V in W . f heißt linear, wenn die Bedingungen*

$$f(x + y) = f(x) + f(y) \quad (4.10)$$

$$f(ax) = af(x), \quad a \in \mathbb{K} \quad (4.11)$$

erfüllt sind.

Wegen (4.11) lassen sich lineare Abbildungen auch durch

$$f(ax + by) = af(x) + bf(y) \quad (4.12)$$

definieren.

Ein Vektorraum ist eine Menge von Elementen, den 'Vektoren', für die eine Addition und eine Multiplikation mit einem Skalar erklärt ist. Die Abbildungen f werden in der folgenden Definition charakterisiert:

Definition 4.4 *Die Bedingungen (4.10) und (4.11) entsprechen den Verknüpfungen von Elementen in einem Vektorraum, sie sind mit den Verknüpfungen im Vektorraum verträglich; die lineare Abbildung f heißt deswegen ein Homomorphismus von Vektorräumen. Darüber hinaus heißt f*

<i>Monomorphismus,</i>	f ist injektiv,
<i>Epimorphismus,</i>	f ist surjektiv,
<i>Isomorphismus,</i>	f ist bijektiv,
<i>Endomorphismus,</i>	$V = W$,
<i>Automorphismus,</i>	f ist bijektiv und $V = W$.

Beispiel 4.1 (Monomorphismus) Es sei V ein Vektorraum über einem Körper $\mathbb{K}$ und $\mathbf{v}_1, \dots, \mathbf{v}_n$ sei ein System von Vektoren aus V . Es werde die Abbildung

$$\mathbf{i} : \mathbb{K}^n \rightarrow V \quad (4.13)$$

betrachtet. $\mathbb{K}^n$ ist die Menge der n -dimensionalen Vektoren, deren Komponenten aus $\mathbb{K}$ stammen, und diese Vektoren werden auf Elemente des Vektorraums V abgebildet:

$$\begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \mapsto x_1 \mathbf{v}_1 + x_2 \mathbf{v}_2 + \dots + x_n \mathbf{v}_n, \quad (4.14)$$

d.h. dem Vektor $(x_1, \dots, x_n)'$ wird der als Linearkombination definierte Vektor $x_1 \mathbf{v}_1 + x_2 \mathbf{v}_2 + \dots + x_n \mathbf{v}_n$ zugeordnet. Die Abbildung $\mathbf{i}$ ist linear. Die Einheitsvektoren $\mathbf{e}_j$ sind als kanonische Basis eingeführt worden (vergl. (2.93) auf Seite 49). Es ist dann

$$\mathbf{i}(\mathbf{e}_j) = \mathbf{v}_j, \quad j = 1, \dots, n \quad (4.15)$$

Die Vektoren $\mathbf{v}_1, \dots, \mathbf{v}_n$ seien linear unabhängig. Die Vektoren $\mathbf{u} = x_1 \mathbf{v}_1 + x_2 \mathbf{v}_2 + \dots + x_n \mathbf{v}_n$ sind dann eindeutig durch den Vektor $(x_1, \dots, x_n)'$ bestimmt, d.h. die Abbildung $\mathbf{i}$ ist injektiv und damit ein Monomorphismus. $\square$

Beispiel 4.2 (Endomorphismus) Ein einfaches Beispiel ist

$$f(a\mathbf{x}) = a\mathbf{x}, \quad a \in \mathbb{R} \quad (4.16)$$

und $\mathbf{x}$ ein n -dimensionaler Vektor. f ist nun eine Abbildung $f : V \rightarrow V$ des Vektorraums $V = \mathbb{R}^n$ in sich selbst. Für den Spezialfall $a = 1$ erhält man

$$\text{id}_x : V \rightarrow V, \quad \mathbf{x} \mapsto \mathbf{x} \quad (4.17)$$

Da V in sich selbst abgebildet wird, ist die Abbildung ein Endomorphismus. $\square$

Satz 4.2 Es seien U, V und W Vektorräume über einem Körper $\mathbb{K}$ und es seien $f : U \rightarrow V$, $g : V \rightarrow W$ lineare Abbildungen. Dann ist die Verknüpfung $g \circ f : U \rightarrow W$ ebenfalls eine lineare Abbildung.

Beweis: Es ist für $a, b \in \mathbb{K}$

$$\begin{aligned}(g \circ f)(a\mathbf{x} + b\mathbf{y}) &= g(f(a\mathbf{x} + b\mathbf{y})) \\ &= g(af(\mathbf{x}) + bf(\mathbf{y})) = ag(f(\mathbf{x})) + bg(f(\mathbf{y})) \\ &= a(g \circ f)(\mathbf{x}) + b(g \circ f)(\mathbf{y}),\end{aligned}$$

wegen der vorausgesetzten Linearität von g und f , – also ist $g \circ f$ ebenfalls linear. $\square$

Beispiel 4.3 (Isomorphismus) Ist f eine bijektive Abbildung, so gibt es für $\mathbf{v} \in V$ genau ein $\mathbf{w} \in W$ mit $f(\mathbf{v}) = \mathbf{w}$ und für alle $\mathbf{v}' \in V$, $\mathbf{v}' \neq \mathbf{v}$ ist $f(\mathbf{v}') \neq \mathbf{w}$. Es existiert eine Abbildung $g : W \rightarrow V$, die dem Element $\mathbf{w}$ genau wieder $\mathbf{v}$ zuordnet; g ist die zu f inverse Abbildung, $g = f^{-1}$. Sind $f : V \rightarrow V'$ und $g : V' \rightarrow V''$ Isomorphismen, so ist auch $(g \circ f) : V \rightarrow V''$ ein Isomorphismus. Denn aus f, g linear und bijektiv folgt, dass auch $g \circ f$ linear und bijektiv ist. Die Eigenschaften der Isomorphie lassen sich wie folgt zusammenfassen:

$$\begin{aligned}(i) \quad & V \simeq V \\ (ii) \quad & V \simeq V' \Rightarrow V' \simeq V \\ (iii) \quad & V \simeq V' \wedge V' \simeq V'' \Rightarrow V \simeq V''\end{aligned}$$

Das "Wesen" eines Isomorphismus ist, dass sich alle Relationen zwischen Objekten in einem Vektorraum V auf den Vektorraum W übertragen. So sei $\mathbf{b}_1, \dots, \mathbf{b}_n$ eine Basis des n -dimensionalen Vektorraums V und es sei $f : V \rightarrow W$ ein Isomorphismus. Die $\mathbf{b}_j$, $j = 1, \dots, n$ werden dann auf die Vektoren $f(\mathbf{b}_1), \dots, f(\mathbf{b}_n) \in W$ abgebildet, und $f(\mathbf{b}_1), \dots, f(\mathbf{b}_n)$ ist eine Basis von W . Es sei weiter $B = (\mathbf{b}_1, \dots, \mathbf{b}_n)$ eine Basis von V . Es sei $\mathbf{x} = (x_1, \dots, x_n)'$ ein Vektor aus $\mathbb{K}^n$, also etwa $x_j \in \mathbb{R}$ und damit $\mathbf{x} \in \mathbb{R}^n$, es sei

$$\begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \mapsto \mathbf{v} = \sum_{j=1}^n x_j \mathbf{b}_j \in V.$$

Dies entspricht einer Abbildung $\mathbf{i}_B : \mathbb{R}^n \rightarrow V$, die ein Isomorphismus von $\mathbb{R}^n$ auf V ist. Denn wenn B eine Basis von V ist, so kann ja jeder Vektor $\mathbf{v} \in V$ eindeutig als Linearkombination der Basisvektoren aus B dargestellt werden, und wegen der Eindeutigkeit ist die Abbildung bijektiv. Damit existiert eine Umkehrabbildung $\mathbf{i}_B^{-1}$

$$\mathbf{i}_B^{-1} : V \rightarrow \mathbb{R}^n, \quad (\text{allgemein } V \rightarrow \mathbb{K}^n) \quad (4.18)$$

$\mathbf{i}_B^{-1}$ heißt auch *Koordinatenabbildung* von V zur Basis B . $\square$

Im folgenden Satz wird eine Bedingung für die Isomorphie einer Abbildung gegeben:

Satz 4.3 *Es seien V und W endlich erzeugte Vektorräume über einem Körper $\mathbb{K}^n$. V und W sind genau dann isomorph, wenn sie dieselbe Dimension besitzen, d.h. wenn*

$$V \simeq W \Leftrightarrow \dim(V) = \dim(W). \quad (4.19)$$

Beweis: Es sei V ein durch eine endliche Basis erzeugter Vektorraum; jeder Vektor aus V wird eindeutig als Linearkombination der Vektoren der Basis erzeugt, d.h. $\text{Bild}_B : \mathbb{K}^n \rightarrow V$ und wegen der Eindeutigkeit existiert die inverse Abbildung $\text{Bild}_B^{-1} : V \rightarrow \mathbb{K}^n$. Damit haben zwei isomorphe Vektorräume auch dieselbe Dimension. Es sei $\dim V = \dim W = n$. Dann folgt $W \simeq \mathbb{K}^n$, $W \simeq \mathbb{K}^n$ und $V \simeq W$. $\square$

Der Satz läßt sich auf unendlich-dimensionale Vektorräume verallgemeinern.

4.3 Kern und Bild einer linearen Abbildung

Der folgende Begriff wird zwar sehr allgemein definiert, hat aber z.B. in der Theorie der linearen Gleichungen eine wichtige Anwendung:

Definition 4.5 *Es sei $f : V \rightarrow W$ eine lineare Abbildung. Es sei*

$$\text{Kern } f = \{\mathbf{v} \in V \mid f(\mathbf{v}) = 0\} \quad (4.20)$$

$$\text{Bild } f = \{w \in W \mid \text{es existiert ein } \mathbf{v} \in V \text{ mit } f(\mathbf{v}) = w\} \quad (4.21)$$

Kern f heißt Kern von f , und Bild f heißt Bild von f . Der Kern von f heißt auch der Defekt von f .

Der Kern einer Abbildung von V in W ist also die Menge der Elemente – der Vektoren – von V , die auf die 0 von W abgebildet werden. Das Bild von f ist die Menge der Vektoren von W , die sich durch die Abbildung f ergeben; natürlich ist $\text{Bild } f = f(V)$.

Beispiel 4.4 Es sei $\mathbf{v} \in \mathbb{R}^3$, d.h. $\mathbf{v}$ sei ein 3-dimensionaler Vektor. Es werde die Abbildung $L : \mathbb{R}^3 \rightarrow \mathbb{R}$, $\mathbf{x} \mapsto \mathbf{v}'\mathbf{x}$, betrachtet, also eine Abbildung, die einem 3-dimensionalen Vektor einen Skalar zuordnet; der Skalar ist als Skalarprodukt des Vektors $\mathbf{v}$ mit einem Vektor $\mathbf{x}$ definiert. Das Bild von L ist die Menge der reellen Zahlen $\mathbb{R}$ für $\mathbf{v} \neq \vec{0}$, und $\text{Bild } L = \{0\}$ für $\mathbf{v} = \vec{0}$. Der Kern von L ist die Menge $\{\mathbf{x} \in \mathbb{R}^3 \mid \mathbf{v}'\mathbf{x} = 0\}$. Aber $\mathbf{v}'\mathbf{x} = 0$ genau dann, wenn $\mathbf{v}$ und $\mathbf{x}$ orthogonal zueinander sind, so dass der Kern aus *allen* Vektoren $\mathbf{x}$ besteht, die auf $\mathbf{v}$ senkrecht stehen, also für $\mathbf{x}_1$ und $\mathbf{x}_2$ aus dem Kern auch alle Linearkombinationen von $\mathbf{x}_1$ und $\mathbf{x}_2$, so dass $\mathbf{v}$ der Normalenvektor einer Ebene ist, vergl. (??), Seite ??). $\square$

Es gilt dann der

Satz 4.4 *Der Kern und das Bild der Abbildung $f : V \rightarrow W$ sind Teilräume von V bzw. W . Insbesondere gilt*

$$\begin{aligned} (i) \quad & f \text{ injektiv} \Leftrightarrow \text{Kern } f = \{0\} \\ (ii) \quad & f \text{ bijektiv} \Leftrightarrow \text{Bild } f = W. \end{aligned} \quad (4.22)$$

Beweis: Es muß nur (i) bewiesen werden, weil (ii) gerade die Definition von Surjektivität ist.

Wegen $f(\vec{0}) = \vec{0}$ folgt $\text{Kern } f \neq \emptyset$, $\text{Bild } f \neq \emptyset$.

Es sei f injektiv, $\mathbf{v}_1, \mathbf{v}_2 \in \text{Kern } f$ und $\lambda \in \mathbb{K}$, so dass $f(\mathbf{v}_1 + \mathbf{v}_2) = f(\mathbf{v}_1) + f(\mathbf{v}_2) = \vec{0} + \vec{0} = \vec{0}$ und $\lambda \mathbf{v}_1 \in \text{Kern } f$, dann $f(\lambda \mathbf{v}_1) = \lambda f(\mathbf{v}_1) = \lambda \cdot 0 = 0$, d.h. $\text{Kern } f$ erfüllt die Bedingungen für einen Vektorraum, ist also ein Teilvektorraum von V . Weiter folgt aus der Injektivität von f , dass für alle $\mathbf{v}_1, \mathbf{v}_2 \in V$, $f(\mathbf{v}_1) = f(\mathbf{v}_2) \Rightarrow \mathbf{v}_1 = \mathbf{v}_2$, so dass $f(\mathbf{v}_1 - \mathbf{v}_2) = f(\vec{0}) = 0 \Leftrightarrow \mathbf{v}_1 - \mathbf{v}_2 \in \text{Kern } f$. Ist also $\text{Kern } f = \{\vec{0}\}$, so folgt, dass f injektiv ist. $\square$

Satz 4.5 *Die Abbildung f sei aus der Menge $\mathcal{L}(U, W)$ der linearen Abbildungen des Vektorraums V in den Vektorraum W . Die Abbildung $f \in \mathcal{L}$ ist injektiv genau dann, wenn $\text{Kern } f = \{\vec{0}\}$.*

Beweis: Es sei $\text{Kern } f = \{\vec{0}\}$. Dann folgt

$$\begin{aligned} f(\mathbf{v}_1) = f(\mathbf{v}_2) &\Rightarrow f(\mathbf{v}_1) - f(\mathbf{v}_2) = 0 \\ &\Rightarrow f(\mathbf{v}_1 - \mathbf{v}_2) = 0 \Rightarrow \mathbf{v}_1 - \mathbf{v}_2 \in \text{Kern } f \Rightarrow \mathbf{v}_1 - \mathbf{v}_2 = 0 \end{aligned}$$

und mithin $\mathbf{v}_1 = \mathbf{v}_2$, also ist f injektiv. Nun sei $\text{Kern } f \neq \{\vec{0}\}$. Dann folgt, dass $f(\mathbf{v}_1) = f(\mathbf{v}_2) = \vec{0}$ für $\mathbf{v}_1 \neq \mathbf{v}_2$, d.h. es werden verschiedene Elemente aus V auf den Nullvektor abgebildet, und mithin ist f nicht injektiv. $\square$

Die Frage nach der Dimensionalität von Vektorräumen spielt in der multivariaten Statistik eine große Rolle. Der folgenden Satz sagt etwas über die Dimensionalität aus:

Satz 4.6 *Es sei V ein n -dimensionaler Vektorraum, $n < \infty$, d.h. $\dim V = n < \infty$, und $f : V \rightarrow W$ sei eine lineare Abbildung von V in einen Vektorraum W . Dann gilt*

$$\dim(\text{Kern } f) + \dim(\text{Bild } f) = n. \quad (4.23)$$

Beweis: Zunächst sei $f = 0$. Dann folgt $\text{Kern } f = V$ (es werden ja *alle* Vektoren bzw. Elemente aus V auf 0 abgebildet), und weiter ist $\text{Bild } f = \{\vec{0}\}$, denn $(\mathbf{v}) = \vec{0} \in W$ für alle $\mathbf{v} \in V$. Also ist $n = 0$.

Nun sei $f \neq 0$ und $\dim(\text{Kern } f) = m$. Dann ist $m < n$, da $\text{Kern } f$ ein Unterraum von V ist. Es sei $m > 0$. Da $\text{Kern } f$ ein Unterraum von V ist existiert eine Basis $(\mathbf{b}_1, \dots, \mathbf{b}_m)$ von Vektoren aus $\text{Kern } f$, so dass $\text{Spann}(\mathbf{b}_1, \dots, \mathbf{b}_m) =$

kern f . Die Basis $(\mathbf{b}_1, \dots, \mathbf{b}_m)$ kann durch Hinzunahme von geeigneten Vektoren $\mathbf{b}_{m+1}, \dots, \mathbf{b}_n$ zu einer Basis von V erweitert werden; für den Fall $m = 0$ kann irgendeine Basis von V gewählt werden.

Die Vektoren $(f(\mathbf{b}_{m+1}), \dots, f(\mathbf{b}_n))$ bilden eine Basis für $\text{Bild } f$. Denn einerseits gilt

$$f(\mathbf{b}_{m+1}), \dots, f(\mathbf{b}_n) \in \text{Bild } f \Rightarrow \text{Spann}(f(\mathbf{b}_{m+1}), \dots, f(\mathbf{b}_n)) \subset \text{Bild } f,$$

andererseits folgt aus $\mathbf{u} \in \text{Bild } f$, $\mathbf{u} = f(\mathbf{v})$ mit $\mathbf{v} = \sum_{j=1}^n c_j \mathbf{b}_j \in V$,

$$\mathbf{u} = \sum_{j=1}^n c_j f(\mathbf{b}_j) = \sum_{j=m+1}^n c_j f(\mathbf{b}_j) \in \text{Spann}(f(\mathbf{b}_{m+1}), \dots, f(\mathbf{b}_n)).$$

Damit gilt

$$\text{Bild } f = \text{Spann}(f(\mathbf{b}_{m+1}), \dots, f(\mathbf{b}_n)).$$

Jetzt ist noch zu zeigen, dass die $\mathbf{b}_{m+1}, \dots, \mathbf{b}_n$ linear unabhängig sind. Dazu sei $\sum_{j=m+1}^n c_j f(\mathbf{b}_j) = \mathbf{0}$. Es sei $\mathbf{u} = \sum_{j=m+1}^n c_j \mathbf{b}_j \in \text{kern } f$, denn $f(\mathbf{u}) = \mathbf{0}$. Ist nun $m = 0$, so ist $c_{m+1} = \dots = c_n = 0$. Ist $m > 0$ so existieren Zahlen $d_1, \dots, d_m$ mit

$$\sum_{j=m+1}^n c_j \mathbf{b}_j = \sum_{j=1}^m d_j \mathbf{b}_j.$$

Aber die Darstellung von Vektoren als Linearkombinationen von Basisvektoren ist eindeutig, so dass $c_{m+1}, \dots, c_n = 0$ folgt. $\square$

Satz 4.7 *Es sei $(\mathbf{v}_1, \dots, \mathbf{v}_n)$ eine Basis des Vektorraums V . Weiter seien $\mathbf{b}_j$, $j = 1, \dots, n$ Vektoren aus einem Vektorraum W . Dann existiert genau eine Abbildung $f : V \rightarrow W$ derart, dass*

$$f(\mathbf{v}_1) = \mathbf{b}_1, \dots, f(\mathbf{v}_n) = \mathbf{b}_n. \quad (4.24)$$

Die Abbildung f ist durch die Angabe der Bildvektoren $f(\mathbf{v}_1), \dots, f(\mathbf{v}_n)$ eindeutig definiert.

Beweis: Zur Existenz von f : Es seien die Vektoren $\mathbf{b}_1, \dots, \mathbf{b}_n \in W$ gegeben. Weiter sei $\mathbf{u} = \sum_{j=1}^n c_j \mathbf{v}_j \in V$, und $f(\mathbf{u}) = \sum_{j=1}^n c_j \mathbf{b}_j \in W$. Dann folgt zunächst, dass f linear ist, denn es sei $\mathbf{w} = \sum_{j=1}^n d_j \mathbf{v}_j$, und

$$f(\mathbf{u} + \mathbf{w}) = \sum_{j=1}^n (c_j + d_j) \mathbf{b}_j = \sum_{j=1}^n c_j \mathbf{b}_j + \sum_{j=1}^n d_j \mathbf{b}_j = f(\mathbf{u}) + f(\mathbf{w}),$$

sowie

$$f(\lambda \mathbf{u}) = \sum_{j=1}^n c_j \lambda \mathbf{b}_j = \lambda \sum_{j=1}^n c_j \mathbf{b}_j = \lambda f(\mathbf{u}),$$

also ist f linear. Die Linearität wiederum impliziert die Eindeutigkeit von f :

$$f(\mathbf{u}) = f\left(\sum_{j=1}^n c_j \mathbf{v}_j\right) = \sum_{j=1}^n c_j f(\mathbf{v}_j) = \sum_{j=1}^n c_j \mathbf{b}_j,$$

d.h. die Koeffizienten c_j , die für die Linearkombination $\mathbf{u}$ anhand der $\mathbf{v}_j$ benötigt werden, sind dieselben, die für die Linearkombination der $\mathbf{b}_j$ zur Darstellung von $f(\mathbf{u})$ anhand der $\mathbf{b}_j$ herangezogen werden müssen. Dies dokumentiert die Eindeutigkeit der Abbildung f (s.a. Satz 4.9, Seite 176). $\square$

Definition 4.6 Der Rang der Abbildung f ist durch

$$\text{rg}(f) = \dim(\text{Bild}f) \quad (4.25)$$

definiert.

Satz 4.8 Es gilt stets

$$\text{rg}(f) \leq \dim W, \quad \text{rg}(f) \leq \dim V, \quad (4.26)$$

und

$$\text{rg}(f) = \dim W \Leftrightarrow f \text{ surjektiv.} \quad (4.27)$$

Beweis: Die Ungleichung $\text{rg}(f) \leq \dim W$ folgt aus dem Satz, dass für $V \subseteq W$ die Ungleichung $\dim V \leq \dim W$ folgt. Es sei V ein n -dimensionaler Vektorraum und $\mathbf{b}_1, \dots, \mathbf{b}_n$ sei eine Basis von V . Es sei $\mathbf{a}_j = f(\mathbf{b}_j)$ für $j = 1, \dots, n$. Dann gilt

$$\dim(\text{Bild}f) = \text{rg}(\mathbf{a}_1, \dots, \mathbf{a}_n),$$

denn $\text{Bild}f = (\mathbf{a}_1, \dots, \mathbf{a}_n)$, mithin folgt $\text{rg}(f) \leq \dim V$. Für $n < \infty$ folgt aus dem Vorangegangenen $\text{rg}(f) = \dim V \Leftrightarrow f$ ist injektiv. $\square$

4.4 Die Matrix einer linearen Abbildung

Multipliziert man eine Matrix M von rechts mit einem geeignet dimensionierten Vektor, so ergibt sich wieder ein Vektor, – in Abschnitt ?? ist dieser Sachverhalt schon beschrieben worden. Ist M eine $(m \times n)$ -Matrix und $\mathbf{x}$ ein n -dimensionaler Vektor, so ist $M\mathbf{x} = \mathbf{y}$ und $\mathbf{y}$ ist ein m -dimensionaler Vektor. Die Multiplikation von M mit $\mathbf{x}$ ordnet dem n -dimensionalen Vektors $\mathbf{x}$ den m -dimensionalen Vektor $\mathbf{y}$ zu, oder: dem Element $\mathbf{x}$ aus dem n -dimensionalen Vektorraum V wird das Element $\mathbf{y}$ aus dem m -dimensionalen Vektorraum W zugeordnet, bzw. $\mathbf{x} \in V$ wird auf $\mathbf{y} \in W$ abgebildet. Man hat generell eine Abbildung $f : V \rightarrow W$, $\mathbf{x} \mapsto \mathbf{y}$. Die Abbildung f wird dabei durch die Matrix M definiert. Die Abbildung ist linear, und man kann sagen, dass M die Matrix der linearen Abbildung f ist.

Satz 4.9 Es sei V ein endlich-dimensionaler Vektorraum über $\mathbb{K}$ mit der Basis $B = (\mathbf{b}_1, \dots, \mathbf{b}_n)$, und W sei ein beliebiger Vektorraum. Zu jedem n -Tupel $(\mathbf{a}_1, \dots, \mathbf{a}_n)$ von Vektoren in W existiert genau eine Abbildung

$$f : V \rightarrow W \text{ mit } f(\mathbf{b}_j) = \mathbf{a}_j, j = 1, \dots, n \quad (4.28)$$

Beweis: Ein beliebiger Vektor $\mathbf{v} \in V$ kann in der Form

$$\mathbf{v} = \sum_{j=1}^n x_j \mathbf{b}_j, \quad x_j \in \mathbb{K}$$

dargestellt werden. Ist f linear, so gilt

$$f(\mathbf{v}) = f\left(\sum_{j=1}^n x_j \mathbf{b}_j\right) = \sum_{j=1}^n f(x_j \mathbf{b}_j) = \sum_{j=1}^n x_j f(\mathbf{b}_j) = \sum_{j=1}^n x_j \mathbf{a}_j.$$

Damit ist f eindeutig bestimmt.

Jetzt ist noch zu zeigen, dass f mit der Eigenschaft (4.28) existiert. Ist $\mathbf{v} = \sum_j x_j \mathbf{b}_j$, und soll $f(\mathbf{v}) = \sum_j x_j \mathbf{a}_j$ sein, so muß f eindeutig sein, da $\mathbf{v}$ eindeutig als Linearkombination definiert ist. f genügt damit (4.19) und f ist linear. $\square$

Beispiel 4.5 Es sei $A = (a_{ij})$ eine beliebige $(m \times n)$ -Matrix über $\mathbb{K}$. Die Spalten von A seien die Vektoren $\mathbf{a}_1, \dots, \mathbf{a}_n$. Nach Satz 4.9 existiert genau eine Abbildung $f : \mathbb{K}^n \rightarrow \mathbb{K}^m$ mit

$$f(\mathbf{e}_j) = \mathbf{a}_j = \begin{pmatrix} a_{1j} \\ a_{2j} \\ \vdots \\ a_{mj} \end{pmatrix}, \quad 1 \leq j \leq n \quad (4.29)$$

wobei $(\mathbf{e}_1, \dots, \mathbf{e}_n)$ die kanonische Basis von $\mathbb{K}^n$ ist. Es sei $\mathbf{x}$ ein beliebiger n -dimensionaler Vektor, der als Linearkombination der kanonischen Basis dargestellt sei, so dass

$$\mathbf{x} = \sum_{j=1}^n x_j \mathbf{e}_j = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}. \quad (4.30)$$

Ein Vektor $\mathbf{u}$ kann dann auch als Linearkombination der Spalten $\mathbf{a}_j$ von A ausgedrückt werden:

$$f(\mathbf{x}) = \mathbf{u} = \sum_{j=1}^n x_j \mathbf{a}_j = \begin{pmatrix} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n \\ \vdots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n \end{pmatrix}. \quad (4.31)$$

$\square$

Diese Charakterisierung einer linearen Abbildung ist allerdings noch ein wenig unvollständig. Man muß sich daran erinnern, dass die Komponenten eines Vektors seine Koordinaten in Bezug auf eine bestimmte Basis des Vektorraums sind, zu dem der Vektor gehört, wie in Definition 2.15 und der darauf folgenden Anmerkung (Seite 47) ausgeführt wird. Jeder Vektor $\mathbf{v} \in V$ läßt sich ja als Linearkombination der Vektoren einer Basis $B = (\mathbf{v}_1, \dots, \mathbf{v}_n)$ von V darstellen:

$$\mathbf{v} = x_1 \mathbf{v}_1 + x_2 \mathbf{v}_2 + \dots + x_n \mathbf{v}_n, \quad (4.32)$$

und die Komponenten x_j des Vektors $\mathbf{x} = (x_1, \dots, x_n)'$ sind dann die Koordinaten von $\mathbf{v}$ in Bezug auf die Basis B . Analog dazu läßt sich jeder Vektor $\mathbf{x} \in W$ als Linearkombination der Vektoren einer Basis $B' = (\mathbf{w}_1, \dots, \mathbf{w}_m)$ von W darstellen:

$$\mathbf{w} = y_1 \mathbf{w}_1 + \dots + y_m \mathbf{w}_m. \quad (4.33)$$

Dann läßt sich insbesondere $f(\mathbf{v}_j) \in W$ als Linearkombination der Basisvektoren in B' darstellen:

$$f(\mathbf{v}_j) = a_{1j} \mathbf{w}_1 + \dots + a_{mj} \mathbf{w}_m = \sum_{i=1}^m a_{ij} \mathbf{w}_i. \quad (4.34)$$

Die Koeffizienten $(a_{1j}, \dots, a_{mj})$ müssen spezifisch für $\mathbf{v}_j$ gewählt werden (weshalb der zweite Index j erscheint) und können in einer Matrix M zusammengefasst werden, $M = A = (a_{ij})$. Wegen (4.32) hat man für einen beliebigen Vektor $\mathbf{v} \in V$

$$f(\mathbf{v}) = x_1 f(\mathbf{v}_1) + x_2 f(\mathbf{v}_2) + \dots + x_n f(\mathbf{v}_n) = \sum_{j=1}^n x_j f(\mathbf{v}_j) = \sum_{j=1}^n x_j \sum_{i=1}^m a_{ij} \mathbf{w}_i,$$

oder

$$f(\mathbf{v}) = \sum_{i=1}^m \left(\sum_{j=1}^n a_{ij} x_j \right) \mathbf{w}_i = \sum_{i=1}^m y_i \mathbf{w}_i \quad (4.35)$$

mit

$$y_i = \sum_{j=1}^n a_{ij} x_j, \quad i = 1, \dots, m \quad (4.36)$$

$f(\mathbf{v})$ wird damit als Linearkombination der Basisvektoren $\mathbf{w}_1, \dots, \mathbf{w}_m$ von W dargestellt, und die Koordinaten (Koeffizienten der $\mathbf{w}_i$) sind durch die y_i gegeben. Man erhält

$$A\mathbf{x} = \mathbf{y} = x_1 \mathbf{a}_1 + x_2 \mathbf{a}_2 + \dots + x_n \mathbf{a}_n, \quad (4.37)$$

wobei die $\mathbf{a}_j$, $j = 1, \dots, n$ die Spaltenvektoren von A sind.

Die Elemente von A werden durch die Wahl der Basis B' für die Vektoren aus W bestimmt, der Vektor $\mathbf{x}$ wird durch die Wahl einer Basis B für die Vektoren aus V bestimmt. Damit hängt $\mathbf{y}$ sowohl von B wie von B' ab. Man schreibt deshalb auch $M_{B'}^B(f)$ für die Matrix A , also $M_{B'}^B(f) = A$.

Illustration: Es sei $f : \mathbb{R}^3 \rightarrow \mathbb{R}^2$, insbesondere sei

$$f \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = A\mathbf{x}. \quad (4.38)$$

Man sieht nun leicht, dass die Spaltenvektoren $\mathbf{a}_j$ von A gerade die Bilder der gewählten Basis von V sind. Dazu werde der Einfachheit halber die kanonische Basis $\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3$ für $\mathbb{R}^3$ angenommen. Setzt man für $\mathbf{x} = (x_1, x_2, x_3)'$ nun der Reihe nach $\mathbf{e}_1, \mathbf{e}_2$ und $\mathbf{e}_3$ ein, so erhält man

$$\begin{aligned} f(\mathbf{e}_1) &= f \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} = A \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} = \begin{pmatrix} a_{11} \\ a_{21} \end{pmatrix} = \mathbf{a}_1 \\ f(\mathbf{e}_2) &= f \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} = A \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} = \begin{pmatrix} a_{12} \\ a_{22} \end{pmatrix} = \mathbf{a}_2 \\ f(\mathbf{e}_3) &= f \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = A \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = \begin{pmatrix} a_{13} \\ a_{23} \end{pmatrix} = \mathbf{a}_3 \end{aligned}$$

Für $\mathbf{x} = x_1\mathbf{e}_1 + x_2\mathbf{e}_2 + x_3\mathbf{e}_3 \in \mathbb{R}^3$ ist $f(\mathbf{x}) = x_1f(\mathbf{e}_1) + x_2f(\mathbf{e}_2) + x_3f(\mathbf{e}_3)$ findet man dann $f(\mathbf{x}) = A\mathbf{x} = x_1\mathbf{a}_1 + x_2\mathbf{a}_2 + x_3\mathbf{a}_3 = \mathbf{y} \in \mathbb{R}^2$.

Hintereinandergeschaltete Abbildungen:

5 Eigenvektoren und Eigenwerte nichtsymmetrischer Matrizen

5.1 Der allgemeine Fall

Der Begriff des Eigenvektors und der des zugehörigen Eigenwerts ergab sich in Abschnitt 3.9.2 bei der Betrachtung einer Koordinatentransformation auf eine natürliche Art und Weise für den Spezialfall symmetrischer Matrizen. Für nichtsymmetrische quadratische Matrizen können ebenfalls Eigenvektoren existieren, die aber nicht notwendig reell sind; so sei A eine orthonormale Matrix. Das Produkt von A mit einem Vektor $\mathbf{x}$ liefert einen Vektor $\mathbf{y}$, der sich von $\mathbf{x}$ möglicherweise durch eine Rotation unterscheidet: so sei

$$A = \begin{pmatrix} \cos \phi & -\sin \phi \\ \sin \phi & \cos \phi \end{pmatrix}.$$

Dann ist

$$A\mathbf{x} = x_1 \begin{pmatrix} \cos \phi \\ \sin \phi \end{pmatrix} + x_2 \begin{pmatrix} -\sin \phi \\ \cos \phi \end{pmatrix} = \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \mathbf{y},$$

und $\mathbf{y}$ ist nur parallel zu $\mathbf{x}$ für diejenigen Werte von ϕ , für die $\cos \phi = 1$ und $\sin \phi = 0$ ist, also z.B. für $\phi = 0$, so dass $A = I$ mit den Spaltenvektoren $(1, 0)'$ und $(0, 1)$. Dies ist der gewissermaßen triviale Fall, bei dem gar keine Rotation erzeugt wird. Man findet allerdings komplexwertige Eigenvektoren mit zugehörigen komplexwertigen Eigenwerten, – für $\phi = \pi/4$ etwa findet man die Eigenvektoren $(i, 1)'$ und $(-i, 1)'$ mit den Eigenwerten $(1+i)/\sqrt{2}$ und $(1-i)/\sqrt{2}$, mit $i = \sqrt{-1}$, wie man durch Nachrechnen bestätigt. Wie komplexe Eigenvektoren und -werte zu deuten sind, wird später noch besprochen werden.

Charakteristische Gleichung einer Matrix: Es sei A eine quadratische Matrix. Aus $A\mathbf{u} = \lambda\mathbf{u}$ folgt $A\mathbf{u} - \lambda\mathbf{u} = \lambda I\mathbf{u} = \vec{0}$, I die Einheitsmatrix; die Dimensionen von I entsprechen denen von A . Diese Gleichung kann in der Form

$$(A - \lambda I)\mathbf{u} = \vec{0} \quad (5.1)$$

geschrieben werden. Diese Gleichung beschreibt ein homogenes Gleichungssystem: In ausgeschriebener Form hat man

$$(a_{11} - \lambda)u_1 + a_{12}u_2 + \cdots + a_{1n}u_n = 0 \quad (5.2)$$

$$a_{21}u_1 + (a_{22} - \lambda)u_2 + \cdots + a_{2n}u_n = 0 \quad (5.3)$$

$$\vdots \quad (5.4)$$

$$a_{n1}u_1 + a_{n2}u_2 + \cdots + (a_{nn} - \lambda)u_n = 0 \quad (5.5)$$

Solche Gleichungssysteme haben nur dann mindestens eine nicht-triviale Lösung (d.h. eine Lösung, die nicht gleich dem Nullvektor $\vec{0}$ ist), wenn die Koeffizientenmatrix nicht vollen Rang hat, d.h. wenn ihre Determinante verschwindet, so dass

$$|A - \lambda I| = \begin{vmatrix} a_{11} - \lambda & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} - \lambda & \cdots & a_{2n} \\ & & \ddots & \\ a_{n1} & a_{n2} & \cdots & a_{nn} - \lambda \end{vmatrix} = 0 \quad (5.6)$$

Entwickelt man die Determinante, so ergibt sich ein Polynom $P(\lambda)$ in λ vom Grad n :

$$|A - \lambda I| = (-1)^n [\lambda^n - \beta_1 \lambda^{n-1} + \beta_2 \lambda^{n-2} + \cdots + (-1)^{n-1} \beta_{n-1} \lambda + (-1)^n \beta_n] = 0. \quad (5.7)$$

Dieser Gleichung bzw. diese Gleichung, wenn man den Faktor $(-1)^n$ wegläßt, heißt *charakteristische Gleichung* der Matrix A , und das Polynom auf der rechten Seite heißt *charakteristisches Polynom* von A . Die Gleichung hat, wie aus der Theorie der Polynome bekannt ist, insgesamt n Lösungen – also mögliche Eigenwerte von A –, von denen aber nicht alle identisch sein müssen. Die Nullstellen von $P(\lambda)$ sind die möglichen Eigenwerte von A . Wie aus dem Fundamentalsatz der Algebra bekannt ist, existieren genau n Nullstellen $\lambda_1, \dots, \lambda_n$, die allerdings nicht alle verschieden sein müssen und die durch komplexe Zahlen $\lambda = \alpha + i\beta$, $i = \sqrt{-1}$, gegeben sein können. Es gilt dabei

Satz 5.1 Ist ein Eigenwert λ der quadratischen Matrix mit reellen Elementen komplex, so existiert ein zweiter Eigenwert $\bar{\lambda}$, der zu λ konjugiert komplex ist, dh gilt $\lambda = \alpha + i\beta$, so ist auch $\bar{\lambda} = \alpha - i\beta$ ein Eigenwert von A .

Beweis: Es ist $(A - \lambda I)\mathbf{u} = 0$. Der Übergang zu konjugiert komplexen Zahlen führt zu $(\bar{A} - \bar{\lambda}I)\bar{\mathbf{u}} = 0$. Aber A ist als reell vorausgesetzt worden, also folgt

$$(A - \bar{\lambda}I)\bar{\mathbf{u}} = 0, \quad (5.8)$$

und dies heißt, dass $\bar{\lambda}$ ebenfalls ein Eigenwert von A ist. $\square$

Links- und Rechtseigenvektoren: Es sei A eine nicht notwendig symmetrische $(n \times n)$ -Matrix, und für einen n -dimensionalen Vektor $\mathbf{u}$ gelte die Beziehung

$$A\mathbf{u} = \lambda\mathbf{u}. \quad (5.9)$$

Dann ist $\mathbf{u}$ ein Eigenvektor von A , und λ ist der zugehörige Eigenwert.

Es sei $B = A'$; gilt

$$B\mathbf{v} = \mu\mathbf{v}, \quad (5.10)$$

so ist $\mathbf{v}$ ein Eigenvektor von B und μ der zugehörige Eigenwert. Es ist

$$(B\mathbf{v})' = \mathbf{v}'B' = \mathbf{v}'A = \mu\mathbf{v}'.$$

$\mathbf{v}$ heißt auch *Linkseigenvektor* von A ; $\mathbf{u}$ in (5.9) heißt dementsprechend *Rechtseigenvektor*. Wegen (5.10) übertragen sich alle Aussagen über Rechtseigenvektoren auf Linkseigenvektoren, was allerdings nicht bedeutet, dass Links- und Rechtseigenvektoren notwendig identisch sind. Notwendig identisch sind sie nur für den Spezialfall symmetrischer Matrizen. Denn wenn $A' = B = A$ gilt, so folgt aus (5.10) $B\mathbf{v} = A\mathbf{v} = \mu\mathbf{v}$, d.h. ein gegebener Linkseigenvektor entspricht einem Rechtseigenvektor. Im Falle $A' \neq A$ gilt der

Satz 5.2 Es sei A eine quadratische, nicht-symmetrische Matrix. Es gelte einerseits $\mathbf{v}'A = \mu\mathbf{v}'$, andererseits $A\mathbf{u} = \lambda\mathbf{u}$ mit $\lambda \neq \mu$. Dann folgt $\mathbf{u}'\mathbf{v} = 0$, d.h. die Links- und Rechtseigenvektoren sind orthogonal zueinander.

Beweis: Multiplikation von $\mathbf{v}'A = \mu\mathbf{v}'$ von rechts mit $\mathbf{u}$ und von $A\mathbf{u} = \lambda\mathbf{u}$ von links mit $\mathbf{v}'$ liefert

$$\begin{aligned} \mathbf{v}'A\mathbf{u} &= \mu\mathbf{v}'\mathbf{u} = \lambda\mathbf{v}'\mathbf{u} \\ \mathbf{v}'A\mathbf{u} &= \lambda\mathbf{v}'\mathbf{u} = \mu\mathbf{v}'\mathbf{u} \end{aligned}$$

Da $\mathbf{v}'A\mathbf{u} - \mathbf{v}'A\mathbf{u} = 0$ folgt $\lambda\mathbf{v}'\mathbf{u} - \mu\mathbf{v}'\mathbf{u} = (\lambda - \mu)\mathbf{v}'\mathbf{u} = 0$, woraus wegen $\lambda - \mu \neq 0$ die Behauptung $\mathbf{v}'\mathbf{u} = 0$ folgt, d.h. $\mathbf{u}$ und $\mathbf{v}$ sind orthogonal. $\square$

Im Falle nicht-symmetrischer Matrizen sind Links- und Rechtseigenvektoren also verschieden, da sie ja orthogonal zueinander sind. Dieses Resultat bedeutet nicht, dass auch die Rechts- und Linkseigenvektoren untereinander orthogonal zueinander sind. Aber die Gültigkeit des folgenden Satzes läßt sich zeigen:

Satz 5.3 *Es sei A eine nicht-symmetrische, quadratische Matrix mit mehr als einem Rechtseigenvektor. Die Rechtseigenvektoren sind linear unabhängig, sofern die zugehörigen Eigenwerte verschieden sind.*

Beweis: Es seien $\mathbf{u}_1$ und $\mathbf{u}_2$ zwei Rechtseigenvektoren von A mit zugehörigen Eigenwerten $\mu \neq \lambda$, für die $\mu \neq \lambda$ gelte. Angenommen, sie seien linear abhängig; dann existieren Koeffizienten a_1 und a_2 ungleich Null derart, dass

$$a_1 \mathbf{u}_1 + a_2 \mathbf{u}_2 = 0 \quad (5.11)$$

Multiplikation von links mit A führt dann auf $a_1 A \mathbf{u}_1 + a_2 A \mathbf{u}_2 = 0$, d.h. auf

$$a_1 \mu \mathbf{u}_1 + a_2 \lambda \mathbf{u}_2 = 0. \quad (5.12)$$

Multipliziert man (5.11) mit λ und subtrahiert (5.11) dann von (5.12), so erhält man

$$a_1(\lambda - \mu) \mathbf{u}_1 + a_2(\lambda - \lambda) \mathbf{u}_2 = 0,$$

woraus wegen $\lambda - \mu \neq 0$ sofort $a_1 = 0$ folgt. Auf analoge Weise folgt $a_2 = 0$, d.h. $\mathbf{u}_1$ und $\mathbf{u}_2$ sind linear unabhängig.

Diese Aussage gilt für irgendzwei Rechtseigenvektoren von A . Hat man also insgesamt drei Eigenvektoren, so sind sie paarweise linear unabhängig, so dass man sagen könnte, sie seien insgesamt linear unabhängig. Das Argument ist aber intuitiv, und ein strenger Beweis ist einer intuitiven Betrachtung stets vorzuziehen. Dieser ergibt sich durch das Prinzip der vollständigen Induktion. Es gebe also $r > 2$ linear unabhängige Eigenvektoren, so dass

$$a_1 \mathbf{u}_1 + a_2 \mathbf{u}_2 + \cdots + a_r \mathbf{u}_r = 0 \text{ genau dann, wenn } a_1 = \cdots = a_r = 0.$$

Es ist zu zeigen, dass dann auch $r + 1$ Eigenvektoren linear unabhängig sind, so dass

$$a_1 \mathbf{v}_1 + a_2 \mathbf{v}_2 + \cdots + a_r \mathbf{v}_r + a_{p+1} \mathbf{v}_{p+1} = 0 \quad (5.13)$$

gilt mit $a_1 = a_2 = \cdots = a_{p+1} = 0$ als einziger Lösung. Da die $\mathbf{u}_j$ Eigenvektoren sind, gilt $A \mathbf{u}_j = \lambda_j \mathbf{u}_j$. Multiplikation von (5.13) mit A führt dann unter Berücksichtigung dieser Beziehung auf die Gleichung

$$a_1 \lambda_1 \mathbf{u}_1 + a_2 \lambda_2 \mathbf{u}_2 + \cdots + a_r \lambda_r \mathbf{u}_r + a_{p+1} \lambda_{p+1} \mathbf{u}_{p+1} = 0. \quad (5.14)$$

Multipliziert man (5.13) mit λ_{p+1} und subtrahiert die Gleichung dann von (5.14), so erhält man

$$a_1(\lambda_1 - \lambda_{p+1}) \mathbf{u}_1 + \cdots + a_p(\lambda_p - \lambda_{p+1}) \mathbf{u}_p = 0,$$

und wegen der vorausgesetzten linearen Unabhängigkeit der $\mathbf{v}_j$, $1 \leq j \leq r$ hat man einerseits $a_1 = \cdots = a_p = 0$ und wegen der ebenso vorausgesetzten Ungleichheit der λ_j folgt dann aus (5.14) $a_{p+1} \lambda_{p+1} \mathbf{u}_{p+1} = 0$. Daraus folgt wegen $\lambda_{p+1} \neq 0$ dann $a_{p+1} = 0$, so dass die $\mathbf{u}_1, \dots, \mathbf{u}_{p+1}$ ebenfalls linear unabhängig sind. $\square$

Der Beweis gilt für eine beliebige quadratische Matrix, also auch für A' und damit für die Rechtseigenvektoren von A' , die aber die Linkseigenvektoren von A sind, so dass deren lineare Unabhängigkeit ebenfalls nachgewiesen ist. Gilt der Spezialfall $A' = A$, ist A also symmetrisch, so folgt sofort, dass in diesem Fall die Linkseigenvektoren gleich den Rechtseigenvektoren sind, und wie bereits gezeigt wurde gilt dann nicht nur die lineare Unabhängigkeit der Eigenvektoren, sondern darüber hinaus auch die Orthogonalität der Eigenvektoren.

Im Folgenden werden nur die Rechtseigenvektoren betrachtet und es wird der Kürze wegen nur von Eigenvektoren geredet; alle Aussagen übertragen sich auf die Linkseigenvektoren. Zunächst soll die Beziehung zwischen einer quadratischen Matrix A und ihren Eigenvektoren und Eigenwerten auf eine andere Art dargestellt werden, die Aufschluß über die Anzahl und Art der Eigenvektoren und -eigenwerte gibt.

Ähnliche Matrizen: Es sei V die Matrix der Eigenvektoren einer beliebigen quadratischen Matrix A . Da die Spaltenvektoren von V linear unabhängig sind, folgt die Existenz der zu V inversen Matrix V^{-1} . Aus $AV = V\Lambda$, Λ die Matrix der zugehörigen Eigenwerte, folgt dann durch Multiplikation von rechts mit V^{-1}

$$A = V\Lambda V^{-1}. \quad (5.15)$$

Durch Multiplikation von rechts mit V und von links mit V^{-1} erhält man hieraus

$$V^{-1}AV = \Lambda. \quad (5.16)$$

Mit diesen beiden Gleichungen hat man einen Spezialfall einer Beziehung zwischen Matrizen, die durch die folgende Definition charakterisiert wird:

Definition 5.1 *Es seien A und B zwei $(n \times n)$ -Matrizen und es existiere eine nichtsinguläre Matrix C derart, dass*

$$B = C^{-1}AC \quad (5.17)$$

gilt. Dann heißen A und B ähnlich.

Offenbar bedeuten (5.15) und (5.16), dass A und Λ ähnlich sind. Da Λ eine Diagonalmatrix ist, heißt A auch *diagonalisierbar*. Damit eine $(n \times n)$ -Matrix diagonalisierbar ist, muß also die Matrix V^{-1} existieren, und diese Matrix existiert, wenn A vollen Rang hat, denn dann hat A n linear unabhängige Eigenvektoren.

Komplexe Eigenwerte und -vektoren: Es ist bisher stets vorausgesetzt worden, dass für eine gegebene quadratische Matrix Eigenwerte und -vektoren existieren. Die Frage ist aber, ob für eine beliebige quadratische Matrix überhaupt Eigenvektoren existieren müssen. Gegeben sei etwa die Matrix

$$A(\phi) = \begin{pmatrix} \cos \phi & -\sin \phi \\ \sin \phi & \cos \phi \end{pmatrix} \quad (5.18)$$

Ein Eigenvektor $\mathbf{v}$ von A muß die Bedingung $A\mathbf{v} = \mathbf{w} = \lambda\mathbf{v}$ erfüllen, d.h. der Vektor $\mathbf{w}$ muß parallel zu $\mathbf{v}$ sein, er darf sich nur in der Länge von $\mathbf{v}$ unterscheiden. Aber für $\phi \neq 0$ und $\phi \neq \pi$ bewirkt A eine Rotation des Vektors $\mathbf{v}$, $\mathbf{w}$ kann also nicht parallel zu $\mathbf{v}$ sein. A hat mit Ausnahme spezieller ϕ -Werte zumindest keinen reellen Eigenvektor. Um die Situation allgemein zu klären, geht man noch einmal auf die Gleichung (5.9) zurück: so dass sich die Eigenwerte von A als Nullstellen des Polynoms ergeben. Speziell für die Matrix (5.18) erhält man

$$|A - \lambda I| = 4\lambda^2 - 4\lambda \cos \phi + 1 = 0; \quad (5.19)$$

auf die Herleitung des Polynoms wird hier verzichtet, da es nur auf die Implikationen von (5.19) ankommt. Man findet die Nullstellen

$$\lambda_1 = \frac{1}{2}(\cos \phi + i \sin \phi), \quad \lambda_2 = \frac{1}{2}(\cos \phi - i \sin \phi), \quad i = \sqrt{-1} \quad (5.20)$$

Die in (5.18) definierte Matrix A hat also zwei komplexe Eigenwerte, reelle Eigenwerte ergeben sich nur für solche ϕ -Werte, für die $\sin \phi = 0$ ist, also etwa für $\phi = 0$, wenn gar keine Rotation der Vektoren stattfindet, oder für $\phi = \pi/2$, wenn eine Rotation um 90° stattfindet.

Es ist also möglich, dass für eine beliebig gewählte quadratische Matrix keine reellen Eigenwerte existieren, dass man aber komplexwertige Eigenwerte finden kann, die als Paare konjugiert komplexer Zahlen auftreten²⁷. Nun hätte man noch gerne die zugehörigen Eigenvektoren bestimmt. Für A findet man zwei:

$$\mathbf{u}_1 = \begin{pmatrix} -i \\ 1 \end{pmatrix}, \quad \mathbf{u}_2 = \begin{pmatrix} i \\ 1 \end{pmatrix}. \quad (5.21)$$

Natürlich ergibt sich die Frage der Deutung von komplexen Eigenwerten und Eigenvektoren. Diese treten etwa bei der Analyse dynamischer Systeme und dementsprechend bei allgemeinen Diskussionen von Zeitreihenproblemen auf; wegen ihrer Bedeutung werden sie im Abschnitt ?? gesondert betrachtet. Hier sollen zunächst noch bestimmte Typen von Matrizen eingeführt werden.

Typen von Matrizen: In der multivariaten Statistik spielen symmetrische Matrizen mit reellen Elementen eine zentrale Rolle, es ist aber trotzdem sinnvoll, auch den allgemeinen Fall einer Matrix mit möglicherweise komplexwertigen Elementen zu betrachten.

Sind die Elemente einer Matrix A komplex, d.h. von der Form $z = x + iy$ mit $i = \sqrt{-1}$, so heißt $\bar{A}$ die zu A konjugierte Matrix; die Elemente von $\bar{A}$ enthalten die hzu z konjugiert komplexen Elemente $\bar{z} = x - iy$. Sind nur die Imaginärteile iy der Elemente einer Matrix A von Null verschieden, so heißt A *imaginär*; in diesem Fall gilt $A = -\bar{A}$. Die Transponierte $\bar{A}'$ einer Matrix A heißt die mit A *assoziierte* Matrix.

²⁷Zwei komplexe Zahlen z und $\bar{z}$ heißen konjugiert komplex, wenn sie sich nur im Vorzeichen des Imaginärteils unterscheiden, wenn also $z = x + iy$ und $\bar{z} = x - iy$ gilt.

Für symmetrische Matrizen gilt $A' = A$, d.h. $a_{ij} = a_{ji}$ für alle i, j . Gilt für eine Matrix die Aussage $a_{ij} = -a_{ji}$, so heißt A *schief-symmetrisch*.

Ein wichtiger Fall ist durch die Gleichung

$$A = \bar{A}' \quad (5.22)$$

definiert; in diesem Fall heißt A *hermitesch*²⁸. Ist $A = \bar{A}$, so sind die Elemente von A alle reell und (5.22) bedeutet einfach, dass A symmetrisch ist. Das der reelle Fall ein Spezialfall ist, gelten alle Aussagen über hermitesche Matrizen auch für reelle symmetrische Matrizen, so dass es Sinn macht, bestimmte Aussagen allgemein für hermitesche Matrizen zu machen.

5.2 Das generalisierte Eigenvektorproblem

Eine Reihe von statistischen Fragestellungen führt auf das *generalisierte Eigenvektorproblem*, so etwa die Frage, ob zwei, an m "Fällen" erhobene Datensätze die gleiche oder eine ähnliche latente Struktur haben oder nicht. So kann man an m Personen (Patienten, etc) Messungen von n Variablen vor und nach einer Intervention (etwa einer Therapie) erheben. Die Frage nach einer Veränderung durch die Intervention (Therapieerfolg) führt auf die Frage, ob sich Vorher- und Nachhermessungen systematisch voneinander unterscheiden. Die Berechnung der Kanonischen Korrelationen kann hier zu Antworten führen. An dieser Stelle soll auf die rein formalen Aspekte derartiger Methoden eingegangen werden.

Definition 5.2 *Es seien A und B symmetrische, positiv semidefinite Matrizen. Dann repräsentiert*

$$A\mathbf{w} = \lambda B\mathbf{w} \quad (5.23)$$

das generalisierte Eigenvektorproblem.

Setzt man $B = I$, I die entsprechende Einheitsmatrix, so reduziert sich das generalisierte Eigenvektorproblem auf das bekannte einfache Eigenvektorproblem.

Generalisierter Rayleigh-Quotient: Der *generalisierte Rayleigh-Quotient* ist durch

$$\rho(\mathbf{w}) = \frac{\mathbf{w}'A\mathbf{w}}{\mathbf{w}'B\mathbf{w}} \quad (5.24)$$

definiert.

Es ist klar, dass es bei diesem Ausdruck nicht auf die Länge der $\mathbf{w}$ ankommt, denn setzt man bei (5.24) $\mu\mathbf{w}$, $\mu \in \mathbb{R}$ für $\mathbf{w}$ ein, so kürzt sich μ sofort heraus. Deswegen kann man für die Länge von $\mathbf{w}$ fordern, dass sie der Bedingung $\mathbf{w}'B\mathbf{w} = 1$ genügt. Man kann dann (5.24) in

$$\rho(\mathbf{w}) = \mathbf{w}'A\mathbf{w}, \quad \mathbf{w}'B\mathbf{w} = 1 \quad (5.25)$$

²⁸Nach dem französischen Mathematiker Charles Hermite (1822 – 1901)

umformulieren.

Die Matrizen A und B sind als symmetrisch vorausgesetzt worden, – aber $B^{-1}A$ ist deswegen nicht notwendig ebenfalls symmetrisch. Dies bedeutet, dass die Folgerungen für den symmetrischen Fall $C\mathbf{w} = \lambda\mathbf{w}$, C eine symmetrische Matrix, nicht mehr zutreffen müssen. So kann man z.B. nicht mehr folgern, dass die verschiedenen Eigenvektoren $\mathbf{w}_j$, die der Gleichung (5.26) genügen, notwendig orthogonal zueinander sein müssen.

Rückführung auf den symmetrischen Fall: B ist als symmetrisch und positiv semidefinit vorausgesetzt worden. Dann kann man die Wurzel $B^{1/2} = P\Lambda^{1/2}P'$ von B bestimmen, – offenbar ist

$$B^{1/2}B^{1/2} = B = P\Lambda^{1/2}P'P\Lambda^{1/2}P' = P\Lambda P',$$

denn $P'P = I$ die Einheitsmatrix. Die Gleichung (5.23) führt durch Multiplikation von links mit B^{-1} auf

$$B^{-1}A\mathbf{w} = \lambda\mathbf{w}. \quad (5.26)$$

Multipliziert man diese Gleichung von links mit $B^{1/2}$, so erhält man

$$B^{1/2}B^{-1}A\mathbf{w} = B^{-1/2}A\mathbf{w} = \lambda B^{1/2}\mathbf{w}.$$

Es sei $\mathbf{v}$ die Transformation von $\mathbf{w}$ mit $B^{-1/2}$, d.h. es sei $\mathbf{w} = B^{-1/2}\mathbf{v}$. Dann erhält man

$$B^{-1/2}AB^{-1/2}\mathbf{v} = \lambda\mathbf{v}. \quad (5.27)$$

Die Matrix $B^{-1/2}AB^{-1/2}$ ist aber symmetrisch: $(B^{-1/2}AB^{-1/2})' = B^{-1/2}AB^{-1/2}$, wegen $A' = A$ und $(B^{1/2})' = B^{1/2}$, und damit hat man mit (5.27) ein Eigenwert- und Eigenvektorproblem der bekannten Art für symmetrische Matrizen. Die Lösungen $\mathbf{v}_j$ sind bekanntlich orthonormal (vergl. Satz ??, Seite ??). Weiter ist

$$\mathbf{v}'A\mathbf{v} = \mathbf{w}'B^{1/2}B^{-1/2}AB^{-1/2}B^{1/2}\mathbf{w} = (B^{1/2}\mathbf{w})'B^{-1/2}(B^{1/2}\mathbf{w})$$

und

$$\mathbf{w}'B\mathbf{w} = \mathbf{w}'B^{1/2}B^{-1/2}BB^{-1/2}B^{1/2}\mathbf{w} = (B^{1/2}\mathbf{w})'(B^{1/2}\mathbf{w}) = \|B^{1/2}\mathbf{w}\|^2,$$

so dass der generalisierte Rayleigh-Quotient die Form

$$\rho = \frac{\mathbf{w}'A\mathbf{w}}{\mathbf{w}'B\mathbf{w}} = \frac{(B^{1/2}\mathbf{w})'B^{-1/2}AB^{-1/2}(B^{1/2}\mathbf{w})}{\|B^{1/2}\mathbf{w}\|^2} \quad (5.28)$$

annimmt. Damit ist der generalisierte Rayleigh-Quotient auf den üblichen Rayleigh-Quotienten für symmetrische Matrizen zurückgeführt. In anderer Formulierung kann man sagen, dass die Lösung für den generalisierten Rayleigh-Quotienten durch die Lösung für den gewöhnlichen Rayleigh-Quotienten in einem *transformierten Raum* gegeben ist. Man kommt damit zu der Aussage (Shaw-Taylor & Christianini (2004), p. 162)²⁹

²⁹Shaw-Taylor, J. Christianini, N.: Kernel Methods for Pattern Analysis. Cambridge University Press, Cambridge 2004

Satz 5.4 Ein beliebiger Vektor $\mathbf{v}$ kann als Linearkombination der $\mathbf{w}_j$, $j = 1, \dots, k$ angeschrieben werden. Für die Eigenvektoren des generalisierten Eigenvektorproblems $A\mathbf{w} = \lambda B\mathbf{w}$ gelten die Relationen

$$\mathbf{w}_i' B \mathbf{w}_j = \delta_{ij}, \quad \mathbf{w}_i' A \mathbf{w}_j = \delta_{ij} \lambda_i, \quad \delta_{ij} = \begin{cases} 1, & i = j \\ 0, & i \neq j \end{cases} \quad (5.29)$$

Beweis: Es war $\mathbf{v} = B^{1/2} \mathbf{w}$, und als Lösungen von (5.27) sind die $\mathbf{v}_j$ orthonormal. Es $i \neq j$ und $\lambda_j \neq 0$. Dann folgt, wegen $A\mathbf{w} = \lambda B\mathbf{w}$ (vergl. (5.26)) und damit $B\mathbf{w} = (1/\lambda)A\mathbf{w}$, folgt

$$0 = \mathbf{v}_i' \mathbf{v}_j = \mathbf{w}_i' B^{1/2} B^{1/2} \mathbf{w}_j = \mathbf{w}_i' B \mathbf{w}_j = \frac{1}{\lambda_j} \mathbf{w}_i' A \mathbf{w}_j;$$

nach (5.23) gilt ja $A\mathbf{w} = \lambda B\mathbf{w}$ und deshalb $(1/\lambda_i)A\mathbf{w} = B\mathbf{w}$. Damit gilt (5.29) für den Fall $i \neq j$.

Nun sei $i = j$; es ist $1 = \mathbf{v}_i' \mathbf{v}_i = \mathbf{w}_i' B^{1/2} B^{1/2} \mathbf{w}_i$, also

$$\lambda_i = \lambda_i \mathbf{v}_i' \mathbf{v}_i = \lambda_i \mathbf{w}_i' B^{1/2} B^{1/2} \mathbf{w}_i = \lambda_i \mathbf{w}_i' B \mathbf{w}_i = \mathbf{w}_i' A \mathbf{w}_i,$$

und dies ist (5.29) für den Fall $i = j$. □

Die Maximierung von (5.29) (Maximierung unter Nebenbedingungen, S. Anhang) führt auf die Gleichung (5.26).

Satz 5.5 Für den generalisierten Rayleigh-Quotienten gilt

$$\rho_1 \leq \rho \leq \rho_2, \quad (5.30)$$

und ρ_1, ρ_2 sind durch die Eigenvektoren definiert, die zum kleinsten bzw. größten Eigenwert korrespondieren.

Der Beweis ergibt sich analog zum Beweis für den Rayleigh-Quotienten für symmetrische Matrizen (Satz von Courant-Fisher, Seite 127).

Satz 5.6 Gilt $A\mathbf{v} = \lambda B\mathbf{v}$ und sind λ und $\mathbf{v}$ die Eigenwerte und Eigenvektoren für den generalisierten Rayleigh-Quotienten, so kann A gemäß

$$A = \sum_{j=1}^r \lambda_j B \mathbf{v}_j (B \mathbf{v}_j)' \quad (5.31)$$

zerlegt werden.

Beweis: Für eine beliebige symmetrische Matrix C mit der Matrix P der Eigenvektoren und der Diagonalmatrix $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_n)$ der Eigenwerte gilt

bekanntlich $C = \sum_j \lambda_j \mathbf{p}_j \mathbf{p}_j'$. Die Matrix $B^{-1/2} A B^{-1/2}$ ist symmetrisch, mithin gilt

$$B^{-1/2} A B^{-1/2} = \sum_{j=1}^r \lambda_j \mathbf{v}_j \mathbf{v}_j'.$$

Multipliziert man von links mit $B^{1/2}$ und von rechts ebenfalls mit $B^{1/2}$, so folgt

$$A = \sum_{j=1}^r \lambda_j B^{1/2} (B^{1/2} \mathbf{v}_j \mathbf{v}_j')' = \sum_j \lambda_j B \mathbf{w}_j (B \mathbf{w}_j)',$$

und das war zu zeigen. $\square$

5.3 Mehrfache Eigenwerte

Es sei A eine $(n \times n)$ -Matrix und es werde die Gleichung $A\mathbf{v} = \lambda\mathbf{v}$ betrachtet: λ ist ein Eigenwert von A und $\mathbf{v}$ der zugehörige Eigenvektor. Es gibt maximal n verschiedene Eigenwerte, d.h. es ist möglich, dass einige Eigenwerte mehrfach vorkommen (multiple Eigenwerte, *multiplicity, repeated eigenvalues*). Ein einfaches Beispiel ist die $(n \times n)$ -Identitätsmatrix I : für jeden n -dimensionalen Vektor $\mathbf{x}$ gilt $I\mathbf{x} = \mathbf{x}$, d.h. jeder Vektor $\mathbf{x}$ ist ein Eigenvektor von I , und alle haben den Eigenwert $\lambda = 1$.

Definition 5.3 Es sei V ein Vektorraum und es sei $V_\lambda = \{\mathbf{v} \in V | A\mathbf{v} = \lambda\mathbf{v}\}$. Dann heißt V_λ der Eigenraum von A zum Eigenwert λ .

Bemerkung: Aus $A\mathbf{v} = \lambda\mathbf{v}$ folgt $(A - \lambda I)\mathbf{v} = \vec{0}$. Diese Gleichung ist ein lineares Gleichungssystem in $\mathbf{v}$, d.h. in den Komponenten von $\mathbf{v}$ als Unbekannten. Bekanntlich heißt die Menge der Vektoren $\mathbf{x}$, die der Gleichung $A\mathbf{x} = \vec{0}$ genügen, der Kern von A : $\text{kern}(A) = \{\mathbf{x} | A\mathbf{x} = \vec{0}\}$. Dementsprechend ist $\text{kern}(A - \lambda I) = V_\lambda$, d.h. der Eigenraum V_λ ist der Kern von $(A - \lambda I)$. $\square$

Da zu jedem Eigenwert λ ein Eigenvektor $\mathbf{v}$ korrespondiert, enthält V_λ zumindest ein Element. Da mit $\mathbf{v}$ auch $a\mathbf{v}$, $a \neq 0$ ein Eigenvektor ist, ist V_λ zumindest ein 1-dimensionaler Teilraum von V . Die Frage ist, ob V_λ stets ein Teilraum von V ist. Man sieht dies leicht ein: sind $\mathbf{v} \neq \mathbf{w}$ aus V_λ , so ist mit $a \in \mathbb{R}$, $b \in \mathbb{R}$ auch $\mathbf{u} = a\mathbf{v} + b\mathbf{w} \in V_\lambda$. Denn wegen $A\mathbf{v} = \lambda\mathbf{v}$, $A\mathbf{w} = \lambda\mathbf{w}$ hat man auch

$$A(a\mathbf{v} + b\mathbf{w}) = aA\mathbf{v} + bA\mathbf{w} = a\lambda\mathbf{v} + b\lambda\mathbf{w} = \lambda\mathbf{u}.$$

Die Eigenwerte von A sind die Nullstellen des Polynoms, das durch die Determinante

$$P_A = |A - \lambda I| = 0$$

definiert ist. Mehrfache Eigenwerte gibt es demnach dann, wenn dieses Polynom mehrfache Nullstellen hat. Man kann nun zeigen, dass, wenn λ eine m -fache Nullstelle von P_A ist, dann die Dimension des Eigenraums V_λ kleiner, höchstens gleich

m ist, d.h. es gibt maximal m linear unabhängige Vektoren in V_λ (der Beweis für diese Aussage wird hier übergangen (vergl. Fischer (1984), Kapitel 4).

Der Begriff des Hauptraums ist eine Verallgemeinerung des Begriffs des Eigenraums:

Definition 5.4 Die Matrix A definiere eine Abbildung f des Vektorraums V in sich selbst, d.h. $f : V \rightarrow V$, und λ sei ein Eigenwert von A (d.h. von f), und $r(\lambda)$ sei die algebraische Vielfachheit von λ . Der Kern der r -fachen Hintereinanderschaltung von $A - \lambda I$ heißt Hauptraum zu λ $H(A, \lambda)$

$$H(A, \lambda) = \{v \in V \mid (A - \lambda I)^r(v) = 0\}. \quad (5.32)$$

Die Elemente von $H(A, \lambda)$ heißen die Hauptvektoren. $v \in V$ ist Hauptvektor der Stufe p , wenn $(A - \lambda I)^p v \neq \vec{0}$.

Anmerkung: Alle Eigenvektoren sind Hauptvektoren der Stufe $p = 1$. □

Der in der folgenden Definition eingeführte Begriff des invarianten Teilraums ist eine weitere Verallgemeinerung des Begriffs des Eigenraums:

Definition 5.5 Die Matrix A definiere eine Abbildung eines Vektorraums in sich selbst: $f : V \rightarrow V$, und es sei $U \subseteq V$. Gilt $f(U) \subseteq U$, d.h. ist die Menge der Vektoren Au wieder eine Teilmenge von U , so heißt U invarianter Teilraum von V , oder einfach f -invariant³⁰.

Anmerkung: Alle Eigenräume sowie alle Haupträume sind invariante Teilräume. □

6 Funktionenräume

6.1 Einführung

Vielfach besteht die Aufgabe einer statistischen Analyse nicht nur darin, bestimmte Vektoren $\mathbf{x} = (x_1, \dots, x_n)'$ "vorherzusagen" oder durch andere Vektoren zu "erklären", sondern stetige Funktionen abzuschätzen. Solche Funktionen können implizit durch Differential- oder Integralgleichungen definiert worden sein, oder sie sollen anhand großer Datenmengen geschätzt werden. Statt Mengen von Vektoren der Art $\mathbf{x}$ betrachtet man deshalb Mengen von Funktionen, insbesondere solche Mengen, die den Bedingungen eines Vektorraumes genügen. Man spricht dementsprechend von Funktionenräumen. Ein Funktionenraum enthält dementsprechend nicht nur Funktionen $f_1, f_2, \dots$, die zum Beispiel über einem Intervall $[a, b] \subset \mathbb{R}$ definiert sind, sondern darüber hinaus auch alle Linearkombinationen

³⁰oder invariant unter f

$\lambda_1 f_1 + \lambda_2 f_2 + \dots$. Hier ist die Anzahl der Funktionen f_j , die in einer Linearkombination auftauchen können, absichtlich nicht auf eine endliche Anzahl $n < \infty$ begrenzt worden. Dies legt nahe, dass Funktionen, analog zu Vektoren der Art $\mathbf{x}$, als Linearkombination von bestimmten Basisfunktionen dargestellt werden können. Die Bestimmung der Bedingungen, unter denen eine solche Repräsentation möglich ist, ist sicherlich relevant für die Frage, wie eine "willkürliche" Funktion, für die also kein expliziter Formelausdruck vorliegt, durch eine Reihenentwicklung repräsentiert werden kann.

Unter Umständen ist eine Darstellung einer Funktion f in der Form

$$f = \sum_{i=1}^{\infty} c_i \phi_i$$

möglich; die c_i sind Koeffizienten von Funktionen ϕ_i , die für die Darstellung von f geeignet sind. Effektiv kann die Summe nur für endliches $n < \infty$ berechnen.

$$f_n = \sum_{i=1}^n c_i \phi_i, \quad f_n \approx f,$$

und die Frage ist dann, ob f_n für größer werdendes n hinreichend schnell gegen die Funktion f konvergiert. Diese Frage führt auf verschiedene Klassen von Funktionen, für die spezifische Bedingungen für $f_n \rightarrow f$ erfüllt sein müssen. Diese Klassen werden durch bestimmte "Räume" gekennzeichnet. Diese Räume werden in den folgenden Abschnitten besprochen.

Die weitere Darstellung fokussiert dann auf Hilberträume, genauer auf die Teilklasse der Hilberträume, die durch die Klasse der über eine Kernfunktion definierbaren Funktionen charakterisiert sind. Die Kernfunktionen spielen insbesondere bei Klassifikationsverfahren, aber auch bei Verallgemeinerungen von Regressions- und PCA-Modellen eine wichtige Rolle.

6.2 Normierte Räume

6.2.1 Definitionen

Der Begriff der Norm ist bereits für Vektoren $\mathbf{x} = (x_1, x_2, \dots, x_n)'$ eingeführt worden:

$$\|\mathbf{x}\| = \left(\sum_{i=1}^n x_i^2 \right)^{1/2}, \quad \|\mathbf{x}\|_p = \left(\sum_{i=1}^n x_i^p \right)^{1/p}, \text{ etc}$$

In analoger Weise ist er für Funktionen definierbar. Zunächst werden die Begriffe der Norm, der Metrik und der Konvergenz allgemein eingeführt; bestimmte Normen werden im Anschluß daran definiert.

Definition 6.1 Es sei $\mathbb{K} = \mathbb{R}$ oder $\mathbb{K} = \mathbb{C}$ und es werde ein Vektorraum X über $\mathbb{K}$ betrachtet. Die Abbildung $\|\cdot\| : X \rightarrow [0, \infty)$ heißt Halbnorm, wenn die Bedingungen

- (a) $\|\lambda x\| = \lambda \|x\|$, $x \in X$; $\lambda \in \mathbb{K}$,
 - (b) $\|x + y\| \leq \|x\| + \|y\|$ für alle $x, y \in X$
- erfüllt sind. Gilt außerdem
- (c) $\|x\| = 0$ dann und nur dann, wenn $x = 0$, so heißt $\|\cdot\|$ Norm.

Anmerkung: X kann ein Funktionenraum sein, d.h. x kann eine Funktion sein. □

Definition 6.2 X sowie wie in Definition 6.1 definiert, $x, y, z \in X$ und $d \in \mathbb{K}$. Es mögen die Bedingungen

- (a) $d(x, y) \geq 0$,
- (b) $d(x, y) = d(y, x)$,
- (c) $d(x, z) \leq d(x, y) + d(y, z)$
- (d) $d(x, y) = 0$ dann und nur dann, wenn $x = y$

gelten. Dann heißt d eine Distanz und (a) bis (d) definieren eine Metrik und X heißt metrischer Raum. Gilt insbesondere $d(x, y) = \|x - y\|$, so induziert die Norm $\|\cdot\|$ die Metrik.

Anmerkung: Mit dem in diesem und ähnlichen Zusammenhängen gebrauchten Wort 'induziert' ist so viel wie 'abgeleitet' gemeint, d.h. die Metrik wird aus der Norm abgeleitet. □

Der Begriff der Konvergenz einer Folge wird ebenfalls aus dem der Norm abgeleitet:

Definition 6.3 Es sei $\{x_n\}_{n \in \mathbb{N}}$ eine Folge von Elementen $x_n \in X$. Die Folge heißt Cauchy-Folge³¹, wenn für alle $\varepsilon > 0$ ein $N = N(\varepsilon) \in \mathbb{N}$ existiert derart, dass für alle $n, m > N(\varepsilon)$, $\|x_n - x_m\| < \varepsilon$. Die Folge $\{x_n\}_{n \in \mathbb{N}}$ heißt konvergent gegen x , wenn für alle $\varepsilon > 0$ die Beziehung $\|x_n - x\| < \varepsilon$ für fast alle n gilt.

Anmerkung: Statt des Ausdrucks 'Cauchy-Folge' wird gelegentlich auch der Ausdruck 'Fundamentalfolge' verwendet. □

Es sei noch einmal daran erinnert, dass die $x \in X$ Zahlen, Vektoren oder Funktionen sein können, – die Begriffsbildungen sind also sehr allgemein.

Definition 6.4 Es werden die folgenden Normen betrachtet:

$$\|x\|_p = \int_G |x(t)|^p dt \quad (p\text{-Norm}) \tag{6.1}$$

$$\|x\|_\infty = \sup_{t \in G} |x(t)| < \infty \quad (\text{Supremumsnorm}) \tag{6.2}$$

³¹Augustin-Louis Cauchy (1789 – 1857), französischer Mathematiker.

Für $p = 2$ erhält man den Spezialfall der Euklidischen Norm.

Für den diskreten Fall hat man die Entsprechungen

$$\|x\|_p = \sum_{i=1}^{\infty} |x_i|, \quad \|x\|_{\infty} = \max_i |x_i|. \quad (6.3)$$

Beispiel 6.1 Ein sehr einfaches Beispiel sind die Folgen

$$x_n = (1, \frac{1}{2}, \frac{1}{3}, \dots, \frac{1}{n}, 0, 0, \dots), \quad x = (1, \frac{1}{2}, \frac{1}{3}, \dots, \frac{1}{n}, \frac{1}{n+1}, \dots).$$

Dann gilt

$$\|x_n - x\|_{\infty} = \frac{1}{n+1} \rightarrow 0 \text{ für } n \rightarrow \infty.$$

□

Den verschiedenen Normen entsprechen Vektorräume von Funktionen:

L^p -Räume: Eine wichtige Klasse von Räumen sind als L^p -Räume definierte Funktionenräume. Dabei steht das L für Lebesgue, nach dem französischen Mathematiker Henri Lebesgue (1875 – 1941), der eine für viele Anwendungen wichtige Integrationstheorie begründete, auf die in diesem Skript allerdings nicht eingegangen werden soll. L^p -Räume sind Funktionenräume

$$L^p = \left\{ x \mid \text{mit der Norm } \|x\|_p = \left(\int_S |x(t)|^p dt \right)^{1/p}, \quad 1 \leq p < \infty \right\}. \quad (6.4)$$

Insbesondere für $p = 2$ erhält man den L^2 -Raum der *quadratintegrierbaren Funktionen*, denn dann soll ja $\int |x(t)|^2 dt < \infty$ gelten. L^2 -Räume spielen in den Anwendungen eine zentrale Rolle.

ℓ^p -Folgen: definierten Norm. Diese Norm ist eine Verallgemeinerung der aus den üblichen Vektorräumen bekannten p -Norm

$$\|\mathbf{x}\|_p = \left(\sum_{i=1}^n x_i^p \right)^{1/p}.$$

Insbesondere ist ℓ^p der Raum der unendlichen Folgen $\{x_n\}_{n=1}^{\infty}$ mit der Norm

$$\|\{x_n\}_{n=1}^{\infty}\|_p = \left(\sum_{i=1}^{\infty} x_i^p \right)^{1/p}, \quad 1 \leq p < \infty \quad (6.5)$$

Den Funktionenräumen entsprechen die Räume von Folgen; in Bezug auf die p -Norm wird von ℓ^p -Folgen gesprochen:

$$\ell^p = \left\{ \{x_n\}_{n \in \mathbb{N}} \mid \text{mit der Norm } \sum_{n=0}^{\infty} |x_n|^p \right\}. \quad (6.6)$$

Der in Definition 6.3 eingeführte Begriff der Cauchy-Folge verdeutlicht, dass die Konvergenz von Folgen an eine bestimmte Norm (die in Definition 6.3 nicht näher spezifiziert wurde) gebunden ist. Ebenso ist klar, dass der Begriff der Metrik (Definition 6.2) mit dem der Norm verknüpft ist.

Ungleichungen: Einige Verallgemeinerungen der bisher eingeführten Distanzen und der mit ihnen verbundenen Metriken ergeben sich aus einer allgemeinen Ungleichung:

Satz 6.1 (Höldersche Ungleichung)³² Es sei $a_k, b_k, p, q \in \mathbb{R}$, $k = 1, \dots, n$, $p > 1$, $q > 1$ und es gelte

$$\frac{1}{p} + \frac{1}{q} = 1.$$

Dann gilt

$$\sum_{k=1}^n |a_k b_k| \leq \left(\sum_{k=1}^n |a_k|^p \right)^{1/p} \left(\sum_{k=1}^n |b_k|^q \right)^{1/q}. \quad (6.7)$$

Beweis: Zur Vorbereitung: es seien $a_j, \lambda_j > 0$, $\sum_{j=1}^n \lambda_j = 1$. Eine Funktion f heißt *konvex*, wenn

$$f(\lambda_1 a_1 + \dots + \lambda_n a_n) \leq f(\lambda_1 a_1) + \dots + f(\lambda_n a_n)$$

gilt. $f(x) = \log x$ ist konvex, also gilt

$$\log(\lambda_1 a_1 + \dots + \lambda_n a_n) \leq \lambda_1 \log a_1 + \dots + \lambda_n \log a_n. \quad (6.8)$$

Dann folgt

$$a_1^{\lambda_1} \dots a_n^{\lambda_n} \leq \lambda_1 a_1 + \dots + \lambda_n a_n, \quad (6.9)$$

denn wegen der Monotonität der $\log$ -Funktion wird man durch Logarithmierung dieser Gleichung sofort auf (6.8) geführt.

Nun sei

$$A := \sum_{k=1}^n |a_k|^p, \quad B := \sum_{k=1}^n |b_k|^q.$$

(6.9) impliziert dann

$$\left(\frac{|a_k|^p}{A} \right)^{1/p} \left(\frac{|b_k|^q}{B} \right)^{1/q} \leq \frac{1}{q} \frac{|a_k|^p}{A} + \frac{1}{q} \frac{|b_k|^q}{B}, \quad k = 1, \dots, n$$

Summiert man über k , so hat man

$$\sum_{k=1}^n \frac{|a_k|^p}{A^{1/p}} \frac{|b_k|^q}{B^{1/q}} \leq \frac{1}{p} \frac{A}{A} + \frac{1}{q} \frac{B}{B} = \frac{1}{p} + \frac{1}{q} = 1,$$

und damit $\sum_{k=1}^n |a_k b_k| \leq A^{1/p} B^{1/q}$. □

³²Otto Ludwig Hölder (1859 – 1937), Mathematiker

Definition 6.5 Es sei M ein metrischer Raum, in dem jede Cauchy-Folge gegen ein $x \in M$ konvergiert. Dann heißt M vollständig. Ein vollständiger, normierter Raum heißt Banachraum³³.

In der Vektoralgebra ist das Skalarprodukt (oder auch inneres Produkt) ohne weitere Einschränkung definiert. Werden statt der Vektoren Funktionen betrachtet, so läßt sich ebenfalls ein Skalarprodukt definieren: sind f und g Funktionen über einem Intervall $[a, b]$, so läßt es sich entsprechend dem inneren Produkt bei n -dimensionalen Vektoren, $\mathbf{x}'\mathbf{y} = \langle \mathbf{x}, \mathbf{y} \rangle = \sum_{i=1}^n x_i y_i$ gemäß

$$\langle f, g \rangle = \int_a^b f(t) \bar{g}(t) dt \quad (6.10)$$

definieren, wobei $\bar{g} = g$ für den Fall, dass die Funktionen reellwertig sind. Es ist aber möglich, dass für die betrachteten Funktionen das innere Produkt $\langle f, g \rangle$ gar nicht existiert. Dies kann der Fall sein, wenn $f, g \in L^p$ mit $p \neq 2$. Denn wenn $|\langle f, g \rangle| < \infty$ existiert, so ist $\langle f, f \rangle = \|f\|^2$, d.h. die Norm ist durch $\|f\| = \sqrt{\langle f, f \rangle}$ gegeben, im Widerspruch zu $\|f\|_p = \left(\int_a^b |f|^p dt \right)^{1/p}$, wenn $f \in L^p$ mit $p \neq 2$. Demnach ist L^2 der einzige Funktionenraum, in dem das "kanonische Skalarprodukt" $\langle f, g \rangle$ erklärt ist.

Das Skalarprodukt hat die folgenden Eigenschaften:

1. $\langle f, g \rangle = \overline{\langle g, f \rangle}$ und $\langle f, g \rangle = \langle g, f \rangle$, wenn f, g reell.
2. $\langle af + bg, h \rangle = a\langle f, g \rangle + b\langle g, h \rangle$,
3. $\langle f, f \rangle \geq 0$ und $\langle f, f \rangle = 0$ für $f = 0$.

Diese Eigenschaften korrespondieren zu denen, die für das Skalarprodukt für n -dimensionale Vektoren bereits gezeigt wurden. Dieser Sachverhalt rechtfertigt die Einführung eines speziellen Typs von Raum, der im folgenden Abschnitt vorgestellt wird.

6.2.2 Anmerkungen zur Konvergenz in Funktionenräumen

Der Ausdruck $f_n \rightarrow f$ signalisiert, dass die Folge $\{f_n\}$ gegen die Funktion f konvergiert. In der "normalen" Analysis wird der Begriff der Konvergenz als Konvergenz von Zahlenfolgen definiert: ist etwa $\{a_n\}$ eine Folge von reellen Zahlen, so bedeutet $a_n \rightarrow a$, dass die a_n mit größer werdendem n immer dichter an der Zahl a liegen. In der Analysis wird zur Klärung dieses Sachverhalts das *Weißerstraßsche Häufungsstellenprinzip* herangezogen:

Häufungsstellenprinzip: Gegeben sei ein endliches Intervall $[a, b]$, in dem unendlich viele Zahlen liegen. Dann existiert in $[a, b]$ mindestens eine Zahl ξ derart,

³³Stefan Banach (1892–1945), polnischer Mathematiker, Begründer der Funktionalanalysis

dass in jedes noch so kleine Intervall um ξ mindestens unendlich viele Zahlen von $[a, b]$ hineinfallen. ξ ist eine *Häufungsstelle*. $\square$

Begründung: Sei der Einfachheit halber $a = 0$, $b = 1$. Eine Zahl $x \in [0, 1]$ ist dann als Dezimalzahl darstellbar. ξ liegt in einem der Teilintervalle

$$[0, .1), [.1, .2), \dots, [.9, 1.0]$$

($[c, d)$ bedeutet, dass c zum Teilintervall gehört, d aber nicht mehr). ξ liege in $[a_j, b_j)$, so dass $\xi = 0.a_1 \dots$ gelten muß. $[a_j, b_j)$ kann nun wieder in zehn Teilintervalle aufgeteilt werden: $.a_1, .a_2, \dots, .a_1, b_j$, und ξ muß in einem dieser Teilintervalle liegen, etwa in dem, das mit a_{1k} beginnt. So kann man weiter fortfahren, so dass man zu der Darstellung $\xi = 0.a_1 a_2 a_3 \dots$ gelangt. Jedes noch so kleine Intervall um ξ enthält immer noch unendlich viele Teilintervalle, die auf die geschilderte Art erzeugt werden können. $\square$

Auf diese Weise bekommt der Konvergenzbegriff, also $|a_n - a| < \varepsilon$ für $n > N(\varepsilon)$, eine klare Bedeutung. Folgen können mehr als einen Häufungspunkt haben. Ein einfaches Beispiel ist die Folge

$$a_{2n-1} = 1 - \frac{1}{n}, \quad a_{2n} = \frac{1}{n}, \quad n = 1, 2, \dots$$

(Courant (1955, p. 57). Für $n \rightarrow \infty$ strebt $1/n$ gegen Null, so dass $a_{2n-1} \rightarrow 1$ und $a_{2n} \rightarrow 0$, d.h. die Folge hat die beiden Häufungspunkte 0 und 1. Weiter kann man die Menge der rationalen Zahlen, d.h. der Zahlen p/q mit $p, q \in \mathbb{N}$ betrachten, die sich als Folge darstellen lassen

$$\frac{1}{1}, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots, \frac{2}{1}, \frac{2}{2}, \frac{2}{3}, \frac{2}{4}, \dots, \frac{3}{1}, \frac{3}{2}, \frac{3}{3}, \frac{3}{4}, \dots, \frac{4}{1}, \frac{4}{2}, \frac{4}{3}, \frac{4}{4}, \dots \quad (6.11)$$

Jede rational Zahl, aber auch jede irrationale³⁴ kann sich als Häufungspunkt darstellen lassen, d.h. es gibt unendlich viele Häufungsstellen. Dann gilt:

Satz 6.2 *In jeder beschränkten unendlichen Zahlenmenge läßt sich eine unendliche Teilfolge $a_1, a_2, a_3, \dots$ herausgreifen, die gegen einen Grenzpunkt ξ konvergiert.*

Beweis: Es sei ξ eine Häufungsstelle aus der betrachteten Menge. Man wählt a_1 derart, dass $|a_1 - \xi| < 1/10$, dann a_2 derart, dass $|a_2 - \xi| < 1/100$, dann a_3 mit $|a_3 - \xi| < 1/1000$, etc. Offenbar strebt a_n mit $n \rightarrow \infty$ gegen ξ . $\square$

Die Aussagen über konvergente Zahlenfolgen übertragen sich nicht ohne weiteres auf Funktionenräume. Dazu werde die auf $-1 \leq x \leq +1$ definierte Funktion

$$f_n(x) = \begin{cases} 1 - n^2 x^2, & x^2 \leq 1/n^2 \\ 0, & x^2 > 1/n^2 \end{cases} \quad (6.12)$$

³⁴Zahlen, die sich nicht als Quotient (= ratio) zweier natürlicher Zahlen darstellen lassen.

Für $x \neq 0$ folgt $\lim_{n \rightarrow \infty} f_n(x) = 0$ und für $x = 0$ folgt $\lim_{n \rightarrow \infty} f_n(x) = 1$. Andererseits läßt sich zeigen, dass $\|f\| = \int f^2 dx = 0$. Offenbar verhält sich eine Folge von bestimmten Funktionen anders als Zahlenfolgen, so dass das Häufungsstellenprinzip nicht unbedingt gilt.

Courant & Hilbert geben zwei mögliche Auswege an: (i) man erweitert den Integral- und Konvergenzbegriff, wie es die Lebesguesche Integraltheorie ermöglicht, (ii) man verengt den Bereich der Funktionen, der betrachtet werden kann. Da die Lebesgue-Theorie hier nicht dargestellt werden soll, bleibt die zweite Möglichkeit. Dazu wird der Begriff der *gleichgradigen Stetigkeit* eingeführt:

Definition 6.6 *Es sei (S, d) ein metrischer Raum (S ein Vektorraum und d eine Distanzfunktion), und $M \subset C(S)$ ($C(S)$ noch nicht definiert!) mit Supremumsnorm. Für M gelte³⁵*

$$\forall \varepsilon > 0 \exists \delta > 0 \forall x \in M, d(s, t) \leq \delta \Rightarrow |x(s) - x(t)| \leq \varepsilon. \quad (6.13)$$

Dann heißt M gleichgradig stetig. (Werner, p. 68)

6.3 Hilberträume

”Weyl, eine Sache müssen Sie mir erklären: Was ist das, ein Hilbertscher Raum? Das habe ich nicht verstanden.”

Diese Frage richtete David Hilbert³⁶ an Hermann Weyl³⁷ nach einem Vortrag (zitiert nach Werner, D.: Funktionalanalysis. p.251). Die Definition des Hilberttraums wird im folgenden Abschnitt gegeben.

6.3.1 Definition und Eigenschaften

Definition 6.7 *Es sei $(X, \|\cdot\|)$ ein normierter Raum, in dem ein Skalarprodukt erklärt ist mit $\langle x, x \rangle^{1/2} = \|x\|$ für alle $x \in X$. Dann heißt X Prähilbertraum. Ist X außerdem vollständig³⁸, so heißt $\mathcal{H}$ Hilbertraum.*

In anderen Worten: ein Hilbertraum $\mathcal{H}$ ist ein vollständiger metrischer L^2 -Raum ($f \in \mathcal{H}$ und $\int |f|^2 dx < \infty$), d.h. ein metrischer Raum mit Skalarprodukt, in dem jede Cauchy-Folge konvergiert. Hilberträume stehen im Zentrum der folgenden Darstellungen.

³⁵D.h. ”Für alle ε größer als Null existiert für alle $x \in M$ ein $\delta > 0$ derart, dass etc

³⁶David Hilbert David Hilbert (1862– 1943), deutscher Mathematiker

³⁷Hermann Weyl (1885–1955), Mathematiker, Physiker, Philosoph, der bei D. Hilbert in Göttingen studiert hatte

³⁸Zur Erinnerung: ein Raum ist vollständig, wenn jede Cauchy-Folge gegen ein Element des Raums konvergiert

In diesem Abschnitt werden einige Eigenschaften von Hilberträumen vorgestellt. Die erste Eigenschaft ist die Gültigkeit der Cauchy-Schwarzschen Ungleichung (auch einfach Schwarzsche Ungleichung):

Satz 6.3 *Es seien $f, g \in L^2$. Dann gilt*

$$|\langle f, g \rangle| \leq \|f\| \cdot \|g\|. \quad (6.14)$$

Beweis: Für den Fall n -dimensionaler Vektorräume ist die Ungleichung bereits bewiesen worden. Es sei $\lambda \in \mathbb{C}$ und $\bar{\lambda}$ sei die zu λ konjugiert komplexe Zahl (der Spezialfall $\lambda \in \mathbb{R}$ ist in der folgenden Argumentation enthalten). Sicherlich gilt

$$\langle x - \lambda y, x - \lambda y \rangle = \|x - \lambda y\|^2 \geq 0.$$

Andererseits ist

$$\begin{aligned} 0 &\leq \langle x - \lambda y, x - \lambda y \rangle = \langle x - \lambda y, x \rangle - \bar{\lambda} \langle x - \lambda y, y \rangle = \\ &\quad \langle x, x \rangle - \lambda \langle y, x \rangle - \bar{\lambda} \langle x, y \rangle + |\lambda|^2 \langle y, y \rangle. \end{aligned}$$

Setzt man insbesondere $\lambda = \langle x, y \rangle \|y\|^{-2} = \overline{\langle x, y \rangle} \|y\|^{-2}$ so folgt $|\langle x, y \rangle|^2 \leq \|x\|^2 \|y\|^2$.

□

Für Funktionen nimmt (6.14) die Form

$$\left| \int_a^b f \bar{g} dt \right|^2 \leq \int_a^b |f|^2 dt \int_a^b |g|^2 dt \quad (6.15)$$

an. Man rechnet leicht nach, dass das Gleichheitszeichen nur dann gilt, wenn $y = \alpha x$, $\alpha \neq 0$. Ebenso gilt

Satz 6.4 *Für $x, y \in X$ gilt die Dreiecksungleichung*

$$\|x + y\| \leq \|x\| + \|y\|. \quad (6.16)$$

Beweis: Es gilt

$$\begin{aligned} \langle x + y, x + y \rangle &= \|x + y\|^2 = \langle x, x \rangle + \langle y, y \rangle + \langle x, y \rangle + \langle y, x \rangle \\ &\leq \|x\|^2 + \|y\|^2 + 2|\langle x, y \rangle|, \end{aligned}$$

denn $\langle x, y \rangle$ kann ja negativ sein. Wegen (6.14) folgt sofort (6.16), da ja $2|\langle x, y \rangle| \geq 0$.

□

Anmerkung: Nach Definition 6.3 ist eine Folge $\{x_m\}_{m \in \mathbb{N}}$ eine Cauchy-Folge, wenn für alle $n, m \in \mathbb{N}$ die Bedingung $\|x_n - x_m\| < \varepsilon$ für alle $\varepsilon > 0$ erfüllt ist. Wegen (6.16) hat man, wenn man x durch $x_n - x$ und y durch $x - x_m$ ersetzt,

$$\|x_n - x_m\| = \|(x_n - x) + (x - x_m)\| \leq \|x_n - x\| + \|x_m - x\|,$$

d.h. jede konvergente Folge ist auch eine Cauchy-Folge. Andererseits muß x kein Element von X sein. Wenn $x \in X$, so ist die Folge *vollständig*. Im Hilbertraum sind alle Folgen vollständig.

Es seien $f, \phi_1, \dots, \phi_n, \dots$ Elemente eines Hilbertraumes $\mathcal{H}$. Dann heißt – wie allgemein in einem metrischen Raum –

$$f = c_1\phi_1 + c_2\phi_2 + \dots + c_n\phi_n + \dots \quad (6.17)$$

eine *Linearkombination*. Gilt

$$c_1\phi_1 + c_2\phi_2 + \dots + c_n\phi_n + \dots = 0$$

dann und nur dann, wenn alle $c_j = 0$, so heißen die ϕ_j *linear unabhängig*. Der metrische Raum heißt *n-dimensional*, wenn er n linear unabhängige Elemente hat und je $n+1$ Elemente des Raums linear abhängig sind. Existiert n nicht, so heißt der Raum *unendlich dimensional*.

Man betrachte die Folgen

$$\begin{aligned} (1, 0, 0, 0, \dots) \\ (0, 1, 0, 0, \dots) \\ (0, 0, 1, 0, \dots) \\ \text{etc} \end{aligned}$$

Diese Folgen sind offenbar linear unabhängig. Daraus folgt, dass der Hilbertsche Folgenraum unendlich ist.

Basisfunktionen: Analog zur Darstellung n -dimensionaler Vektoren kann man die Funktionen ϕ_j in (6.17) als *Basisfunktionen* betrachten, wobei wie bei den n -dimensionalen Vektoren orthogonale Basisfunktionen von besonderem Interesse sind. Dies sind Funktionen, die die Bedingung

$$\langle \phi_i, \phi_j \rangle = \delta_{ij} = \begin{cases} 0, & i \neq j \\ 1, & i = j. \end{cases} \quad (6.18)$$

erfüllen, wobei δ_{ij} wieder das Kronecker-delta ist.

Definition 6.8 *Es seien $\phi_1, \phi_2, \phi_3, \dots$ orthonormale Funktionen; sie bilden ein orthonormales Basissystem .*

Bilden die ϕ_j in (6.17) ein orthonormales Basissystem, so folgt

$$\langle \phi_j, f \rangle = c_j, \quad (6.19)$$

denn

$$\langle \phi_k, f \rangle = \sum_j c_j \langle \phi_k, \phi_j \rangle = c_j,$$

da ja $\langle \phi_k, \phi_j \rangle = 0$ für $j \neq k$ und $\langle \phi_k, \phi_j \rangle = 1$ für $j = k$. Es gilt nun der

Satz 6.5 Es sei $\{x_j\}$ eine Folge linear unabhängiger Elemente aus $\mathcal{H}$. Dann läßt sich die Folge in eine Folge orthogonaler Funktionen $\{\phi_j\}$ transformieren.

Beweis: Es werde $\phi_1 = x_1/\|x_1\|$ gesetzt; dann ist $\|\phi_1\| = 1$. Es sei $z_2 = x_2 - \langle x_2, \phi_1 \rangle \phi_1$. Offenbar sind z_2 und ϕ_1 orthogonal, denn

$$\langle \phi_1, z_2 \rangle = \langle \phi_1, x_2 \rangle - \langle x_2, \phi_1 \rangle \langle \phi_1, \phi_1 \rangle = \langle \phi_1, x_2 \rangle - \langle \phi_1, x_2 \rangle = 0,$$

wegen $\langle \phi_1, \phi_1 \rangle = 1$. z_2 wird normalisiert: $\phi_2 = z_2/\|z_2\|$. In dieser Weise kann fortgefahren werden, bis man die orthonormale Basis $\phi_1, \phi_2, \dots$ gewonnen hat; dies ist das Schmidtsche Orthogonalisierungsverfahren³⁹. $\square$

Besselsche Ungleichung:

Satz 6.6 Es sei $\{\phi_j\}_{j \in \mathbb{N}}$ und $x \in \mathcal{H}$, und es gelte die Reihenentwicklung $x = \sum_{j=1}^{\infty} c_j \phi_j$. Dann gilt die Besselsche Ungleichung⁴⁰

$$\sum_{j=1}^{\infty} |c_j|^2 = \sum_{j=1}^{\infty} |\langle \phi_j, x \rangle|^2 \leq \|x\|^2. \quad (6.20)$$

Beweis: Es sei zunächst an den Satz des Pythagoras erinnert; es seien $x, y, z \in \mathcal{H}$ und es gelte (i) $z = x + y$, (ii) $\langle x, y \rangle = 0$, d.h. x und y seien orthogonal. Dann gilt

$$\|z\|^2 = \|x\|^2 + \|y\|^2. \quad (6.21)$$

Denn $\|z\|^2 = (x + y)'(x + y) = \|x\|^2 + \|y\|^2 + 2\langle x, y \rangle = \|x\|^2 + \|y\|^2$ wegen $\langle x, y \rangle = 0$. Dies läßt sich auf beliebige Summen $z = x_1 + x_2 + \dots + x_n + \dots$ verallgemeinern, wenn $x_i \perp x_j$ für $i \neq j$, da dann ja alle Skalarprodukte $\langle x_i, x_j \rangle$ verschwinden. Nun sei für beliebiges $N \in \mathbb{N}$

$$x_N = x - \sum_{j=1}^N \langle \phi_j, x \rangle \phi_j.$$

Wegen der Orthogonalität der ϕ_j folgt dann

$$\|x\|^2 = \|x_N + \sum_{j=1}^N \langle \phi_j, x \rangle \phi_j\|^2 = \|x_N\|^2 + \left\| \sum_{j=1}^N \langle \phi_j, x \rangle \phi_j \right\|^2 \geq \left\| \sum_{j=1}^N \langle \phi_j, x \rangle \phi_j \right\|^2.$$

Da N beliebig gewählt werden kann, folgt (6.20). $\square$

Kleinste-Quadrate: Es sei $f \in \mathcal{H}$ eine Funktion aus einem Hilbertraum und es existiere eine Orthonormalbasis $\{\phi_j\}_{j \in \mathbb{N}}$, so dass $f = \sum_{j=1}^{\infty} c_j \phi_j$. Gesucht ist eine Abschätzung $\hat{f} = \sum_{j=1}^N a_j \phi_j$ für $N < \infty$, wobei die a_j geeignet zu wählende Koeffizienten sind. Es zeigt sich, dass die KQ-Schätzungen für die a_j den c_j entsprechen:

³⁹Erhard Schmidt (1876 – 1958), deutscher Mathematiker

⁴⁰Friedrich Wilhelm Bessel (1784 – 1846), Astronom, Mathematiker, Geodät und Physiker.

Satz 6.7 Die KQ-Schätzungen für die a_j sind verzerrungsfreie Schätzungen für die Entwicklungskoeffizienten c_j von f bezüglich der Funktionen ϕ_j .

Beweis: Es wird die Norm

$$Q_n(a_1, a_2, \dots) = \|f - \sum_{j=1}^n a_j \phi_j\|^2 \quad (6.22)$$

betrachtet; die a_j sollen so gewählt werden, dass Q_n minimal wird. Es ist

$$(f - \sum_{j=1}^n a_j \phi_j)^2 = f^2 + (\sum_{j=1}^n a_j \phi_j)^2 - 2f \sum_{j=1}^n a_j \phi_j,$$

und das Integral darüber ist

$$M := \int f^2 dx + \sum_{j=1}^n a_j^2 \int \phi_j^2 dx + 2 \sum_{j \neq k} a_j a_k \int \phi_j \phi_k dx - 2 \sum_{j=1}^n a_j \int f \phi_j dx,$$

woraus wegen $\int f \phi_j dx = c_j$ und $a_j^2 - 2a_j c_j = (a_j - c_j)^2 - c_j^2$

$$M = \|f\|^2 + \sum_{j=1}^n (a_j - c_j)^2 - \sum_{j=1}^n c_j^2 \quad (6.23)$$

folgt. Da alle Terme auf der rechten Seite größer oder gleich Null sind, wird M offenbar minimal, wenn $\sum_{j=1}^n (a_j - c_j)^2 = 0$, d.h. wenn

$$a_j = c_j \quad \text{für alle } j \quad (6.24)$$

□

Die Güte der Approximation wird natürlich von der Anzahl n der Terme in der Summe abhängen. So kann es gelingen, durch geeignete Wahl von n das Minimum von M unter eine beliebig kleine Schranke zu senken. Wenn dies möglich ist, so heißt die Menge der Funktionen $\phi_1, \phi_2, \dots$ ein *vollständiges orthogonales Funktionensystem* und es gilt die

Vollständigkeitsrelation

$$\sum_{j=1}^{\infty} c_j^2 = \|f\|^2. \quad (6.25)$$

Sie charakterisiert offenbar einen Spezialfall der Besselschen Ungleichung (6.20) und kann verallgemeinert werden. Es seien f und g zwei Funktionen, die im selben Basissystem dargestellt werden können, so dass

$$f = \sum_j c_j \phi_j, \quad g = \sum_j d_j \phi_j$$

mit

$$c_j = \langle f, \phi_j \rangle, \quad d_j = \langle g, \phi_j \rangle.$$

Dann folgt

$$\sum_{j=1}^{\infty} c_j d_j = \langle f, g \rangle, \quad (6.26)$$

denn

$$\|f + g\|^2 = \|f\|^2 + \|g\|^2 + 2\langle f, g \rangle = \sum_{j=1}^{\infty} (c_j + d_j)^2 = \sum_{j=1}^{\infty} (c_j^2 + d_j^2 + 2c_j d_j),$$

woraus sofort (6.26) folgt. Das Bemerkenswerte an der Beziehung (6.26) ist, dass das Skalar- oder innere Produkt von f und g gleich dem Skalarprodukt der Entwicklungskoeffizienten von f und g ist.

Hinreichend für die Vollständigkeit ist, dass die Funktionen stetig sind. Dass die Funktionen in Reihen entwickelt werden können ergibt sich aus der gleichmäßigen Konvergenz der Summe $\sum_{j=1}^{\infty} c_j d_j$ ⁴¹.

Definition 6.9 *Es sei $\mathcal{H}$ ein Hilbertraum und $T \subset \mathcal{H}$ sei eine Teilmenge von $\mathcal{H}$. T heißt abgeschlossen in $\mathcal{H}$, wenn es zu jedem Element $x \in \mathcal{H}$ und für alle $\varepsilon > 0$ eine Linearkombination*

$$u = c_1 z_1 + c_2 z_2 + \cdots + c_n z_n, \quad z_j \in T, \quad j = 1, \dots, n$$

gibt derart, dass $\|u - x\| < \varepsilon$ gilt. Teilmengen T mit dieser Eigenschaft heißen Grundmengen.

Es sei $\{\phi_j\}$ ein Orthonormalsystem für einen Hilbertraum $\mathcal{H}$; die Frage ist, ob auch jeder Vektor $h \in \mathcal{H}$ als Linearkombination der ϕ_j dargestellt werden kann. Da die ϕ_j eine Teilmenge T bilden, heißt die Frage also, ob die $\{\phi_j\}$ eine Grundmenge bilden. Darüber gibt der folgende Satz Auskunft.

Satz 6.8 *Es sei $\{\phi_j\}$ ein Orthonormalsystem eines Hilbertraumes $\mathcal{H}$. $\{\phi_j\}$ ist abgeschlossen genau dann, für jeden Vektor $f \in \mathcal{H}$ die Beziehung*

$$\|f\|^2 = \sum_{j=1}^{\infty} |\langle f, \phi_j \rangle|^2 \quad (6.27)$$

gilt.

⁴¹Es sei $\{f_n\}$ eine Folge von Funktionen; sie konvergiert gleichmäßig gegen eine Funktion f , wenn $\lim_{n \rightarrow \infty} |f_n(x) - f(x)|$ für alle $x \in G$, G der Definitionsbereich der f_n und f .

Beweis: Dies ist offenbar die Beziehung (6.25). Tatsächlich ist (6.27) notwendig, damit für endliches n

$$\|f - \sum_{j=1}^n c_j \phi_j\|^2 < \varepsilon, \quad n > N(\varepsilon)$$

gilt. (6.27) ist auch hinreichend. Denn es sei

$$f = \sum_{j=1}^n \|c_j \phi_j\|^2 + z_n.$$

Aus der Orthonormalität folgt dann

$$\|f\|^2 = \sum_{j=1}^n |c_j|^2 + \|z_n\|^2,$$

(6.23) impliziert dann (6.25) $\|z_n\|^2 = \|f\|^2 - \sum_{j=1}^n |c_j|^2 = 0$, d.h. $f = \lim_{n \rightarrow \infty} \sum_{j=1}^n c_j \phi_j$. □

Definition 6.10 Ein Hilbertraum $\mathcal{H}$ heißt separierbar, wenn $\mathcal{H}$ eine abzählbare Grundmenge hat.

Satz 6.9 Ein Hilbertraum $\mathcal{H}$ enthält genau dann eine vollständige Folge von orthonormalen Basisfunktionen, wenn $\mathcal{H}$ separierbar ist.

Beweis: Gegeben sei eine abzählbare Menge von Elementen $g_1, g_2, \dots$ einer Grundmenge. Streicht man darin die linear abhängigen Elemente, so erhält man eine Folge $u_1, u_2, \dots$ von linear unabhängigen Elementen, die orthogonalisiert werden können. Die resultierende Menge $\phi_1, \phi_2, \dots$ von orthonormalen Elementen bilden wieder eine Grundmenge und sind demnach abgeschlossen und vollständig.

Ist andererseits $\phi_1, \phi_2, \dots$ ein vollständiges und abgeschlossenes System von orthogonalen Elementen, so ist eine abzählbare Grundmenge, so dass $\mathcal{H}$ separierbar ist. □

Satz 6.10 In jedem separierbarem Hilbertraum $\mathcal{H}$ existieren vollständige Orthonormalsysteme $\{\phi_j\}$, durch die jedes Element von $f \in \mathcal{H}$ eindeutig dargestellt werden kann.

Beweis: Es ist schon gezeigt worden, dass $f \in \mathcal{H}$ in der Form

$$f = \sum_{j=1}^{\infty} c_j \phi_j, \quad \text{mit } c_j = \langle f, \phi_j \rangle, \quad \sum_{j=1}^{\infty} |c_j|^2 = \|f\|^2. \quad (6.28)$$

dargestellt werden kann. Es muß noch die Eindeutigkeit der Darstellung gezeigt werden. Allgemein gilt⁴²

$$\lim_{n \rightarrow \infty} \langle f_n, g \rangle = \langle \lim_{n \rightarrow \infty} f_n, g \rangle, \quad (6.29)$$

denn nach der Schwarzschen Ungleichung hat man

$$|\langle f_n, g \rangle - \langle f, g \rangle| = |\langle f_n - f, g \rangle| \leq \|f_n - f\| \cdot \|g\|,$$

so dass $\langle f_n, g \rangle$ gegen $\langle f, g \rangle$ konvergiert. Es sei

$$f = \lim_{n \rightarrow \infty} s_n = \lim_{n \rightarrow \infty} \sum_{j=1}^{\infty} c_j \phi_j,$$

dann folgt nach (6.29)

$$\langle f, \phi_j \rangle = \langle \lim_{n \rightarrow \infty} s_n, \phi_j \rangle = \lim_{n \rightarrow \infty} \langle s_n, \phi_j \rangle = c_j.$$

□

6.3.2 Lineare Operatoren

Elemente eines Hilbertraums $\mathcal{H}$ sind Vektoren, insbesondere auch Funktionen, – der Begriff des Vektors umfaßt also nicht nur Objekte wie $\mathbf{x} = (x_1, \dots, x_n)'$, sondern eben auch Funktionen, die auf einem bestimmten Intervall definiert sind. Man kann Abbildungen auf Hilberträumen definieren. Wird dabei einem Element $f \in \mathcal{H}$ eine Zahl $a \in \mathbb{C}$ (als Spezialfall $a \in \mathbb{R}$) zugeordnet, so heißt die Abbildung *Funktional*. Wird dem Element f ein anderes Element $g \in \mathcal{H}$ zugeordnet, so heißt die Abbildung *Operator*. Funktionen $f, g, \dots$ sind Abbildungen, und die Ausdrücke Funktional und Operator sollen helfen, Verwechslungen mit den Abbildungen $f, g, \dots$ und Abbildungen zwischen Vektorräumen, bei denen Funktionen eine Zahl oder eine andere Funktion zugeordnet wird zu vermeiden.

Definition 6.11 *Es sei T ein Operator in einem Hilbertraum, für den die folgenden Bedingungen gelten:*

1. $T(f_1 + f_2) = Tf_1 + Tf_2$, (*Additivität*)
 2. $T(af) = aTf$, $a \in \mathbb{C}$, (*Homogenität*)
 3. *Es existiert eine Konstante M und $\|Tf\| \leq M \cdot \|f\|$ (*Beschränktheit*).*
- für alle $f \in \mathcal{H}$. Dann heißt der Operator linear.

Anmerkung: Der Operator heißt *stetig*, wenn aus $f_n \rightarrow f$ stets $Tf_n \rightarrow Tf$ folgt. Es läßt sich zeigen, dass 3. Stetigkeit impliziert. □

⁴²Meschkowski (1962), p. 18

Definition 6.12 Es sei T ein linearer Operator. Es sei $\|T\|$ die kleinste der Schranken M in Definition 6.11, 3., so dass $\|Tf\| \leq \|T\| \cdot \|f\|$. Dann heißt $\|T\|$ die Norm des Operators.

Beispiele für einen Operator sind (1) die Multiplikation mit einem Element $g \in \mathcal{H}$, d.h. $Tf = f \cdot g$, (2) $Tf(x) = df(x)/dx = f'(x)$, also die Differentiation des Elements f , wobei allerdings Definition 6.11, 3. nicht erfüllt ist. Von besonderem Interesse sind die *Integraloperatoren*

$$g(u) = Tf(u) = \int_a^b k(x, u) f(x) dx, \quad (6.30)$$

wobei $k(x, u)$ eine *Kernfunktion* ist. Auf Kernfunktionen wird in Abschnitt 6.3.3 ausführlich eingegangen.

Es gelten die folgenden Aussagen:

$$\begin{aligned} (cT)f &= c \cdot Tf, & c \in \mathbb{C} \\ (T_1 + T_2)f &= T_1f + T_2f \\ (T_1 \cdot T_2)f &= T_1(T_2f), & f \in \mathcal{H} \end{aligned}$$

und für die Normen gilt

$$\|cT\| = |c|\|T\|, \quad \|T_1 + T_2\| \leq \|T_1\| + \|T_2\|, \quad \|T_1T_2\| \leq \|T_1\|\|T_2\|.$$

Diese Beziehungen implizieren sofort

$$\|T^n\| \leq \|T\|^n.$$

Ein Operator ordnet einem Element $f \in \mathcal{H}$ ein anderes Element $h \in \mathcal{H}$ zu. Dementsprechend kann man das innere Produkt $\langle h, g \rangle = \langle Tf, g \rangle$ bilden; analog dazu kann man das innere Produkt $\langle f, Tg \rangle$ bilden, wobei allerdings nicht $\langle Tf, g \rangle = \langle f, Tg \rangle$ gelten muß. Es könnte aber einen anderen Operator T^* geben, so dass $\langle Tf, g \rangle = \langle f, T^*g \rangle$ gilt; man kann dann davon ausgehen, dass zwischen den Operatoren T und T^* eine Beziehung besteht. Dazu wird zunächst die folgende Definition eingeführt:

Definition 6.13 Es sei $\mathcal{H}$ ein Hilbertraum und T sowie T^* seien zwei Operatoren, für die die Beziehung

$$\langle Tf, g \rangle = \langle f, T^*g \rangle \quad (6.31)$$

gelte. Dann heißt T^* der zu T adjungierte Operator.

Anmerkung: Es läßt sich zeigen, dass, wenn T beschränkt ist, auch T^* beschränkt ist; dann gilt $\|T\| = \|T^*\|$. $\square$

Operatoren lassen sich durch Matrizen darstellen:

Satz 6.11 *Es sei $\mathcal{H}$ ein separierbarer Hilbert-Raum. Ein linearer Operator T in $\mathcal{H}$ ist durch die unendliche Matrix*

$$M(T) = \begin{pmatrix} a_{11} & a_{12} & a_{13} & \cdots \\ a_{21} & a_{22} & a_{23} & \cdots \\ a_{31} & a_{32} & a_{33} & \cdots \\ \vdots & \vdots & \vdots & \vdots \end{pmatrix} \quad (6.32)$$

definiert, wobei

$$a_{ik} = (Ty_k, y_i), \quad \sum_{i=1}^{\infty} |a_{ik}|^2 < \infty, \quad k = 1, 2, 3, \dots \quad (6.33)$$

und $\{y_i\}$ ist ein vollständiges Orthonormalsystem in $\mathcal{H}$, und

$$\left| \sum_{i=1}^p \sum_{k=1}^q a_{ki} \alpha_i \bar{\beta}_k \right| \leq M \cdot \sqrt{\sum_{i=1}^p |\alpha_i|^2} \sqrt{\sum_{k=1}^q |\beta_k|^2}, \quad (6.34)$$

$\alpha_i, \beta_k \in \mathbb{C}$, $0 < M \in \mathbb{R}$.

Beweis: Es sei $f = \sum_{k=1}^{\infty} \alpha_k y_k$ gegeben. Nach (6.33) folgt

$$Tf = \sum_{k=1}^{\infty} \beta_k y_k,$$

wobei β_k noch unbekannt ist. Aber es muß, unter Benutzung von $T \lim s_n = \lim Ts_n$ mit $s_n = \sum_{k=1}^{\infty} \alpha_i y_i$

$$\beta_k = (Tf, y_k) = \left(T \sum_{k=1}^{\infty} \alpha_i y_i, y_k \right) = \sum_{k=1}^{\infty} (\alpha_i; Ty_i, y_k) = \alpha_i a_{ki} \quad (6.35)$$

gelten, wegen

$$\|Ts_n - Tf\| \leq \|T\| \|s_n - f\|.$$

Für $f = y_i$ folgt aus (6.33)

$$Ty_i = \sum_{k=1}^{\infty} a_{ki} y_k,$$

woraus wiederum (6.34) folgt. und (6.35) folgt wegen

$$|\langle Tf, g \rangle| \leq \|Tf\| \|g\| \leq \|T\| \|f\| \|g\|$$

und $f = \sum_{k=1}^{\infty} \alpha_i y_i$ und $g = \sum_{k=1}^{\infty} \beta_k y_k$, $\|Tf\| = M$. □

Definition 6.14 *Eine Teilmenge A eines metrischen Raumes X ist genau dann relativ kompakt, falls jede Folge in A eine in X konvergente Teilfolge hat.*

Definition 6.15 E und F seien Banachräume und K sei eine lineare Abbildung $K: E \rightarrow F$ von E in F . K heißt kompakter Operator, wenn eine der folgenden äquivalenten Eigenschaften erfüllt ist:

1. Der Operator K bildet jede beschränkte Teilmenge von E auf eine relativ kompakte Teilmenge von F ab.
2. Das Bild der offenen (oder der abgeschlossenen) Einheitskugel in E ist relativ kompakt in F .
3. Jede beschränkte Folge $\{x_n\}$ in E besitzt eine Teilfolge $\{x_{n_k}\}$, sodass $K x_{n_k}$ in F konvergiert.

Definition 6.16 Ist f die betrachtete Funktion, so soll für alle $\varepsilon > 0$ ein von f abhängiges $\delta = \delta(\varepsilon)$ existieren derart, dass für $|x_1 - x_2| < \delta(\varepsilon)$ die Beziehung $|f(x_1) - f(x_2)| < \varepsilon$ gilt. Es läßt sich zeigen, dass alle f mit $\int_a^b |df/dx|^2 dx \leq M$, M eine fixe Konstante, gleichgradig stetig sind.

Eigenwerte und Eigenfunktionen eines Operators:

Definition 6.17 Es sei $\mathcal{H}$ ein Hilbertraum, T ein Operator und es gelte

$$Tf = \lambda f. \quad (6.36)$$

Dann heißt f eine Eigenfunktion und λ ein Eigenwert des Operators T .

Ein für das Folgende wichtiges Beispiel sind die Eigenfunktionen und Eigenwerte des Integraloperators (6.30); man hat dann

$$\lambda f(u) = \int_a^b k(x, u) f(x) dx. \quad \lambda \neq 0 \quad (6.37)$$

Definition 6.18 Der Operator T heißt symmetrisch, wenn $T = T^*$, d.h. wenn T selbstadjungiert ist und wenn $\mathcal{D} = \mathcal{H}$.

Dann gilt der

Satz 6.12 Der Operator T sei symmetrisch. Existieren Eigenwerte für T , so sind sie reell und die zugehörigen Eigenvektoren bzw Eigenfunktionen sind orthogonal.

Beweis: Für einen symmetrischen Operator ist $\langle Tf, f \rangle$ ist stets reell, denn

$$\langle Tf, f \rangle = \langle f, T^* f \rangle = \langle f, Tf \rangle = \overline{\langle T, f \rangle}.$$

Weiter ist $\langle Tf, f \rangle = \lambda \langle f, f \rangle$, woraus folgt, dass λ reell ist. Nun seien f, g Eigenfunktionen zu verschiedenen Eigenwerten λ, μ . Dann gilt

$$\langle Tf, g \rangle = \langle \lambda f, g \rangle = \lambda \langle f, g \rangle,$$

und

$$\langle f, Tg \rangle = \langle f, \mu g \rangle = \mu \langle f, g \rangle,$$

und da T symmetrisch ist folgt $\lambda \langle f, g \rangle = \mu \langle f, g \rangle$. Wegen der Voraussetzung $\lambda \neq \mu$ folgt dann $\langle f, g \rangle = 0$, d.h. f und g sind orthogonal. $\square$

Definition 6.19 *Es sei T ein Operator in einem Hilbertraum $\mathcal{H}$. T heißt vollstetig, wenn aus jeder in der Norm beschränkten Folge $\{f_n\}$ von Elementen aus $\mathcal{H}$ ($\|f_n\| < C$) eine Teilfolge $\{f_{n'}\}$ ausgewählt werden kann, für die $Tf_{n'}$ konvergent ist.*

Satz 6.13 *Jeder symmetrische vollstetige Operator hat mindestens einen und höchstens abzählbar viele Eigenwerte. Zu jedem Eigenwert gehören nur endlich viele Eigenvektoren (Eigenfunktionen), und die zu verschiedenen Eigenwerten gehörenden Eigenvektoren (Eigenfunktionen) sind orthogonal.*

Beweis: Meschkowski (1963), p. 33. $\square$

Spektralsatz für kompakte Operatoren auf Hilberträumen (Werner, p. 270)

Der Spektralsatz ist damit das Analogon zur Hauptachsentransformation der linearen Algebra. (Elaborieren!)

Satz 6.14 *(Spektralsatz für kompakte Operatoren) Gegeben sei ein Hilbertraum $\mathcal{H}$ und T ein selbstadjungierter Operator. Dann existiert ein Orthonormalsystem $e_1, e_2, \dots$ sowie eine Nullfolge $\lambda_1, \lambda_2, \dots$, $\lambda_j \neq 0$, so dass*

$$Tv = \sum_k \lambda_k \langle v, e_k \rangle e_k, \quad \forall v \in \mathcal{H}, \quad (6.38)$$

wobei die λ_k die von Null verschiedenen Eigenwerte von T , und e_k ist der zugehörige Eigenvektor (bzw die zugehörige Eigenfunktion; es gilt $\|T\| = \sup_k \lambda_k$).

Anmerkung: Das Orthonormalsystem kann zu einer Orthonormalbasis B erweitert werden. Dazu muß dann die Orthonormalbasis von $\text{kern}T$ hinzugenommen werden, – und $\text{kern}T$ kann unendlichdimensional sein. Man hat dann

$$v = \sum_k \langle v, e_k \rangle e_k \Rightarrow Tv = \sum_{v \in B} \lambda_e \langle v, e \rangle e. \quad (6.39)$$

6.3.3 Kernfunktionen

Gelegentlich werden Funktionen implizit durch Gleichungen des Typs

$$f(s) = \phi(s) - \int_G K(x, t) \phi(t) dt \quad (6.40)$$

definiert; Gleichungen dieser Art heißen Integralgleichungen bzw. Integraloperatoren

$$Tf(s) = \int_G K(s, t)f(t)dt. \quad (6.41)$$

Die hierin auftretende Funktion $K(s, t)$ heißt *Kernfunktion*. Für $G = [a, b]$ ist die Integralgleichung vom *Fredholm-Typ*, für $G = [0, s]$ sind sie vom *Volterra-Typ*. Im Allgemeinen werden Lösungen vom Typ

$$\lambda x - Tx = y, \quad \lambda \in \rho(T) \quad (6.42)$$

gesucht. Dabei ist $\rho(T)$ die Resolventenmenge:

Definition 6.20 Es sei $T \in L(X)$, d.h. T ein Element der Menge der Abbildungen $\{f|f : X \rightarrow X\}$. Dann heißt

$$\rho(T) = \{\lambda \in \mathbb{K} | (\lambda - T)^{-1} \text{ existiert in } L(X)\}. \quad (6.43)$$

die Resolventenmenge von T .

Zur Erläuterung: Nach (6.42) ist $\lambda x - Tx = (\lambda - T)x = y$, und wenn $(\lambda - T)^{-1}$ existiert, so hat man $x = (\lambda - T)^{-1}y$. Der Kern von $(\lambda - T)$ ist dann nicht leer, d.h. die durch $(\lambda - T)$ definierte Abbildung ist nicht injektiv. Die Menge der λ , für die $\lambda - T$ nicht injektiv ist, heißt auch das *Punktspektrum* von T . Das Punktspektrum entspricht den von Null verschiedenen Eigenwerten einer Matrix in der linearen Algebra. Es sei S eine Orthonormalbasis von $\text{kern } T$ und λ_k, e_k wie in Satz 6.14 gewählt werden, d.h. die von Null verschiedenen Eigenwerte und die dazu korrespondierenden Eigenvektoren bzw Eigenfunktionen von T . Dann gilt $(\lambda - T)x = y$ genau dann, wenn

$$\lambda \sum_k \langle x, e_k \rangle e_k + \lambda \sum_{e \in S} \langle x, e \rangle e - \sum_k \lambda_k \langle x, e_k \rangle e_k = \sum_k \langle y, e_k \rangle e_k + \sum_{e \in S} \langle y, e \rangle e,$$

und dies ist äquivalent zu

$$\begin{aligned} \langle x, e \rangle &= \frac{1}{\lambda} \langle y, e \rangle, \quad \forall e \in S \\ \langle x, e_k \rangle &= \frac{1}{\lambda - \lambda_k} \langle y, e_k \rangle, \quad \forall k \in \mathbb{N} \end{aligned} \quad (6.44)$$

und die Lösung ergibt sich in der Form

$$x = \sum_k \frac{1}{\lambda - \lambda_k} \langle y, e_k \rangle e_k + \frac{1}{\lambda} \sum_{e \in S} \langle y, e \rangle e. \quad (6.45)$$

Integralgleichungen lassen sich auch über eine Reihenentwicklung der Kernfunktion lösen, die ebenfalls in eine Reihenentwicklung von f münden kann. Derartige Reihenentwicklungen machen von orthonormalen Funktionensystemen Gebrauch; über die sich Kernfunktionen gemäß

$$K(x, y) = \sum_{j=1}^{\infty} \phi_j(x) \overline{\phi_j(y)} \quad (6.46)$$

definieren lassen; es ist möglich, dass die ϕ_j komplexe Funktionen sind, und $\overline{\phi_j(y)}$ ist die zu ϕ_j konjugiert komplexe Funktion. Dazu sei gleich angemerkt, dass nicht für jedes System von orthonormalen Funktionen eine Summe der Form (6.46) existiert. Ein Beispiel sind die Funktionen $\phi_j(x) = \sin(jx)/\sqrt{\pi}$; sie bilden ein orthonormales Funktionensystem, aber die Summe $K(x, y)$ wie in (6.46) existiert nicht (vergl. Meschkowski (1963), p. 6): es sei etwa $x = y$. Dann folgt

$$K(x, x) = \sum_{j=1}^{\infty} \phi_j^2(x) = \frac{1}{\pi} \sum_{j=1}^{\infty} \sin^2(jx),$$

und das Konvergenzkriterium

$$\forall x \neq, \forall \varepsilon > 0 \exists N = N(\varepsilon), \|\sin^2(jx) - \sin^2(kx)\| < \varepsilon, \quad j \neq k, j, k > N$$

ist nicht erfüllt.

Satz 6.15 (*Reproduziernde Kerne*) Die Funktion f sei als konvergente Reihe

$$f(x) = \sum_{j=1}^{\infty} a_j \phi_j(x), \quad \sum_{j=1}^{\infty} |a_j|^2 < \infty \quad (6.47)$$

darstellbar, und K sei eine Kernfunktion, die in der Form (6.46) darstellbar sei. Dann gilt

$$f(u) = \langle f(x), K(x, \bar{u}) \rangle = \int_a^b f(x) \overline{K(x, \bar{u})} dx. \quad (6.48)$$

Beweis: Die ϕ_j bilden nach Voraussetzung ein Orthonormalsystem, so dass

$$\langle \phi_j(x), \phi_k(x) \rangle = \delta_{jk}, \quad (6.49)$$

δ_{jk} das Kronecker-delta. In das Skalarprodukt $\langle f(x), K(x, \bar{u}) \rangle$ setzt man die Entwicklung (6.47) ein:

$$\begin{aligned} \langle f(x), K(x, \bar{u}) \rangle &= \left\langle \sum_{j=1}^{\infty} a_j \phi_j(x) \overline{\phi_j(u)} \right\rangle \\ &= \sum_{j=1}^{\infty} \sum_{k=1}^{\infty} \langle a_j \phi_j(x), \phi_k(x) \overline{\phi_k(u)} \rangle = \sum_{j=1}^{\infty} a_j \phi_j(u) = f(u) \end{aligned}$$

□

Aronszajn (1950) hat eine allgemeine Theorie der Kernfunktionen formuliert und wählte dazu (6.48) in der Form

$$f(y) = \langle f(x), K(x, y) \rangle_x, \quad (6.50)$$

als Ausgangspunkt (Meschkowski (1962), p. 42), wobei der Index x andeutet, in Bezug auf welche Variable das Skalarprodukt zu bestimmen ist. (6.50) gestattet es, eine Kernfunktion zu definieren ohne vorher ein Orthonormalsystem zu spezifizieren. Dazu wird zunächst der Begriff des reproduzierenden Kerns definiert:

Definition 6.21 *Es sei $\mathcal{E}$ eine Klasse von Funktionen, in der ein Hilbertraum $\mathcal{H}$ definiert sei. Die Funktion $K(x, y)$ mit $x, y \in \mathcal{E}$ heißt reproduzierender Kern, wenn die Bedingungen*

1. $K(x, y)$ gehört für jedes y als Funktion von x zu $\mathcal{H}$,
2. $K(x, y)$ erfüllt (6.50) (die reproduzierende Eigenschaft).

Satz 6.16 *Es gelten die Relationen*

$$K(x, y) \geq 0, \quad K(x, y) = \overline{K(y, x)} \quad (6.51)$$

$$|K(x, y)|^2 \leq K(x, x) \cdot K(y, y). \quad (6.52)$$

Beweis: Um (6.51) zu zeigen wendet man (6.50) auf die Kernfunktion K selbst an, denn für fixes y ist $K(x, y)$ als Funktion von x ja ein Element von $\mathcal{H}$. Dementsprechend hat man

$$K(y, y) = \langle K(x, y), K(x, y) \rangle = \|K(x, y)\|^2 \geq 0,$$

womit die erste Aussage von (6.51) gezeigt ist. Weiter gilt

$$\langle K(u, y), K(u, x) \rangle = K(x, x), \quad \langle K(u, x), K(u, y) \rangle = K(y, x),$$

und hieraus folgt die zweite Aussage. Die Beziehung (6.52) ergibt sich sofort aus der Schwarzschen Ungleichung. $\square$

Folgerung: Für beliebige $\lambda_j \in \mathbb{C}$ erhält man aus (6.51)

$$\sum_{j=1}^n \sum_{k=1}^n \lambda_j \bar{\lambda}_k K(x_j, y_k) \geq 0, \quad x_j \in \mathcal{E}. \quad (6.53)$$

Denn wegen (6.50) hat man

$$0 \leq \left\langle \sum_{j=1}^n \lambda_j K(x, x_j), \sum_{k=1}^n \lambda_k K(x, x_k) \right\rangle = \sum_{j=1}^n \sum_{k=1}^n \lambda_j \bar{\lambda}_k K(x_j, x_k).$$

$\square$

Satz 6.17 *Es sei $\mathcal{H}$ ein Hilbertraum. Dann existiert höchstens ein reproduzierender Kern für $\mathcal{H}$,*

Beweis: Es werde angenommen, dass es zwei Kernfunktionen K und K' gibt. Dann gilt aber

$$\begin{aligned}\|K(x, y) - K'(x, y)\|^2 &= \langle K - K', K - K' \rangle_x - \langle \langle K - K', K' \rangle_x \langle K - K', K \rangle_x \\ &= K(y, y) - K'(y, y) - K(y, y) + K'(y, y) = 0,\end{aligned}$$

d.h. $K = K'$. □

Wenn es möglich ist, dass ein Hilbertraum keinen reproduzierenden Kern hat, so ist es von Interesse zu wissen, unter welchen Bedingungen er einen Kern hat (man erinnere sich: $x, y \in \mathcal{H}$ sind Funktionen, $f(x), f(y)$ sind Funktionale).

Satz 6.18 *Es sei $\mathcal{H}$ ein Hilbertraum mit Funktionen $f(y)$, $y \in \mathcal{E}$. $\mathcal{H}$ hat einen reproduzierenden Kern $K(x, y)$, wenn für alle $y \in \mathcal{E}$ und alle $f \in \mathcal{H}$ gilt, dass f ein lineares Funktional von $\mathcal{H}$ ist.*

Beweis: $\mathcal{H}$ habe einen Kern K . Aus (6.50) und der Schwarzschen Ungleichung folgt

$$\begin{aligned}|f(y)| &= |\langle f(y), K(x, y) \rangle| \leq \|f\| \cdot \|K(x, y)\| = \|f\| \langle K(x, y), K(x, y) \rangle^{1/2} \\ &= \|f\| K(y, y)^{1/2}\end{aligned}$$

woraus folgt, dass f in der Tat beschränkt ist. Man hat

$$\frac{|f(y)|}{\|f\|} \leq K(y, y)^{1/2}, \quad (6.54)$$

woraus die Eigenschaften linearer Funktionale folgen.

Es sei umgekehrt f ein lineares Funktional. Dann existiert ein von y unabhängige Funktion $g_y(x) \in \mathcal{H}$ derart, dass

$$f(y) = \langle f(x), g_y(x) \rangle$$

(Satz von Riesz). Dann hat $g_y(x)$ die Eigenschaft der Reproduktion und ist demnach eine Kernfunktion, d.h. $g_y(x, y) = K(x, y)$. □

Satz 6.19 *Es sei $\mathcal{H}$ ein Hilbertraum mit Kernfunktion $K(x, y)$; für $f \in \mathcal{H}$, die an der Stelle y den Wert 1 annehmen, gilt*

$$\frac{\|K(x, y)\|}{K(y, y)} \leq \frac{\|f(x)\|}{|f(y)|}, \quad (6.55)$$

Für alle $f \in \mathcal{H}$ mit $\|f\| = 1$ ist $|f(y)|$ mit $y \in \mathcal{E}$ höchstens gleich $\|K(x, y)\| K(x, y)\|^{-1}$ für $x = y$.

Beweis: Die Aussage folgt aus (6.54). $\square$

Satz 6.19 impliziert, dass Hilberträume mit Kernfunktion stets in der Form

$$f(x) = \sum_{j=1}^{\infty} a_j \phi_j(x) \quad (6.56)$$

dargestellt werden können. Ist der Hilbertraum überdies separierbar und hat ein vollständiges Orthonormalsystem $\{\phi_j\}$, so ergibt sich die Darstellung

$$f(x) = \sum_{j=1}^{\infty} \langle f, \phi_j \rangle \phi_j(x), \quad (6.57)$$

wobei die Reihe gleichmäßig in jeder Teilmenge $\mathcal{E}' \subset \mathcal{E}$ konvergiert, in der $K(x, x)$ gleichmäßig beschränkt ist. (s. Meschkowski 1962, p. 48).

Man gelangt schließlich zum

Satz 6.20 Satz von Mercer *Es sei $K(x, y)$ eine auf $[0, 1] \times [0, 1]$ definierte, stetige (Kern-)Funktion und $T_k : L^2[0, 1] \rightarrow L^2[0, 1]$ sei der zugehörige Integraloperator. Weiter gelte $K(s, t) = \overline{K(t, s)}$ für alle $s, t \in [0, 1]$, d.h. T sei selbstadjungiert. Weiter seien $\lambda_1, \lambda_2, \dots$ die Eigenwerte von T_k mit den zugehörigen Eigenfunktionen $e_1, e_2, \dots$. Für $T > 0$ gilt dann*

$$K(s, t) = \sum_{j=1}^{\infty} \lambda_j e_j(s) \overline{e_j(t)}, \quad \forall s, t \in [0, 1], \quad (6.58)$$

und die Konvergenz ist absolut und gleichmäßig.

Beweis: Vorbereitung 1.: man braucht den Satz(VI.4.3), p. 279 Voraussetzung S.v. Mercer, $x \in [0, 1]$,

$$(T_k x)(s) = \sum_{j=1}^{\infty} \lambda_j \langle x, e_j \rangle e_j, \quad \forall s \in [0, 1]$$

und Konvergenz ist absolut und glm.

Denn: T_k ist ein kompakter Operator auf einem Prähilbertraum $[0, 1]$, Konvergenz in L^2 . muß nur glm Konvergenz gezeigt werden s. p. 279:

$$\begin{aligned} \sum_{j=1}^{\infty} |\lambda_j e_j(s)|^2 &= \sum_{j=1}^{\infty} |(T_k e_j)(s)|^2 \\ &= \sum_{j=1}^{\infty} |\langle k(s, \cdot), \bar{e}_j \rangle|^2 \\ &\leq \|k(s, \cdot)\|_{L^2}^2 \text{ Bessel} \\ &= \int_0^1 |k(s, t)|^2 dt \\ &\leq \|k\|_{\infty}^2. \end{aligned}$$

Für $\varepsilon > 0$, $\in \mathbb{N}$, mit

$$\sum_{j=N}^{\infty} |\langle x, e_j \rangle|^2 \leq \varepsilon^2.$$

Dann noch Cauchy-Schwarz

$$\sum_{j=N}^{\infty} |\lambda_j \langle x, e_j \rangle e_j(s)| \leq \left(\sum_{j=N}^{\infty} |\lambda_j e_j(s)|^2 \right)^{1/2} \left(\sum_{j=N}^{\infty} |\langle x, e_j \rangle|^2 \right)^{1/2} \leq \|k\|_{\infty} \varepsilon,$$

und damit ist die glm Konvergenz bewiesen.

Für einen selbstadjungierten Operator $T \in K(\mathcal{H})$ sind die beiden Aussagen äquivalent:

- (i) Alle Eigenwerte von T sind ≥ 0 ,
- (ii) T ist positiv. Denn (i) $\Rightarrow$ (ii) folgt aus Spektralsatz, und (ii) $\Rightarrow$ (i) ist wegen $\langle Tx, x \rangle = \lambda \|x\|^2$ Eigenvektor klar.

Weiter gilt $k(t, t) \geq 0$, $\forall t \in [0, 1]$. Dann noch Satz von Dini: T sei kompakter metrischer Raum, $f, f_1, f_2, \dots, T \rightarrow \mathbb{R}$ stetig und $f_1 \leq f_2 \leq \dots$ und $f = \sup f_n$ punktweise. Dann konvergiert $\{f_n\}$ glm gegen f .

Jetzt kommt erst der eigentliche Beweis für Mercer.

Es sei

$$K_n(s, t) = \sum_{j=1}^n \lambda_j e_j(s) \overline{e_j(t)}.$$

Nach dem Spektralsatz gilt

$$\begin{aligned} \langle T_{k-k_n} x, x \rangle &= \langle T_k x, x \rangle \langle T_{k_n} x, x \rangle \\ &= \sum_{j>n} \langle x, e_j \rangle \langle e_j, x \rangle \\ &= \sum_{j>n} \lambda_j |\langle x, e_j \rangle|^2 \geq 0 \end{aligned}$$

da alle $\lambda_j \geq 0$. Da $K(t, t) \geq 0$ für alle $t \in [0, 1]$. Also

$$K(t, t) - k_n(t, t) \geq 0$$

(s. a. Courant & Hilbert), und damit

$$\sum_{j=1}^n \lambda_j |e_j(t)|^2 \leq K(t, t) \leq \|K\|_{\infty}.$$

Deswegen konvergiert die Reihe links. Nach Cauchy-Schwarz hat man

$$\begin{aligned} \sum_{j=1}^{\infty} \lambda_j e_j(s) \overline{e_j(t)} &= \sum_{j=1}^{\infty} \sqrt{\lambda_j} e_j(s) \sqrt{\lambda_j} \overline{e_j(t)} \\ &\leq \left(\sum_{j=1}^{\infty} \lambda_j |e_j(s)|^2 \right)^{1/2} \left(\sum_{j=1}^{\infty} |e_j(t)|^2 \right)^{1/2} \\ &\leq \|K\|_{\infty} \end{aligned}$$

und $\sum_{j=1}^{\infty} \lambda_j e_j(s) \overline{e_j(t)}$ konvergiert absolut für alle $s, t \in [0, 1]$. Jetzt noch glm Konvergenz (Satz von Dini) (absolute und glm Konvergenz). $\square$

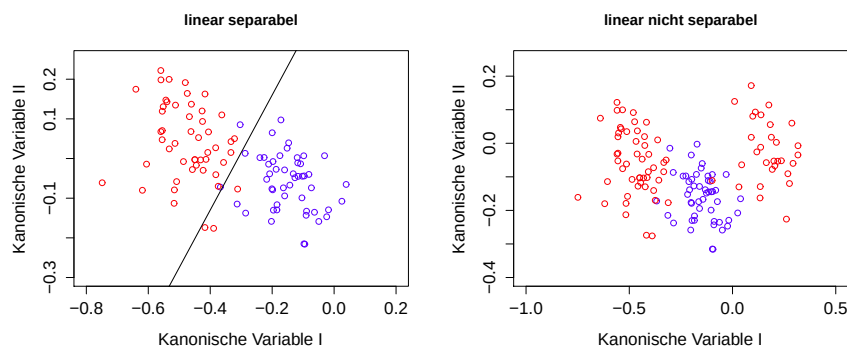
7 Kernmethoden

Vorbemerkung: Satz 2.2 = Parameter steht senkrecht auf Hyperebene.

7.1 Klassifikation

In klassischen Linearen Diskriminanzanalyse von Fisher wird versucht, Hyperebenen in einem n -dimensionalen Prädiktorraum zu finden, die verschiedene Gruppen optimal voneinander trennen. Abbildung 13 (a) zeigt ein Beispiel von zwei Gruppen in einem 2-dimensionalen Raum, bei die Gruppen durch eine Hyperebene – hier eine Gerade – separiert werden können. Die Abbildung (b) zeigt ein Beispiel für eine nicht-trennbare Konfiguration von zwei Gruppen. Die Gerade in (a) wird

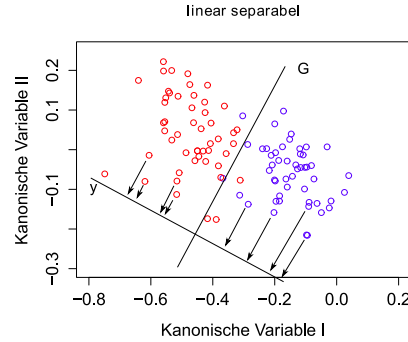
Abbildung 13: Linear trennbare und linear nicht trennbare Konfigurationen



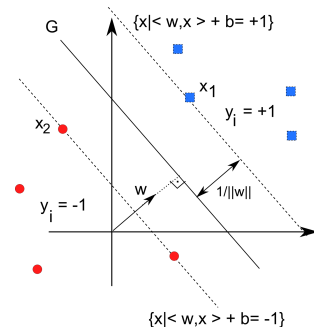
so bestimmt, dass die Projektion der Punkte auf eine zur Hyperebene orthogonale Gerade möglichst keine, in jedem Fall eine nur minimale Überlappung der Projektionen minimal wird, – dadurch wird auch die Hyperebene selbst optimal

(vergl. Abb. 14). Der Punkt bei diesem Vorgehen ist, dass *alle* Punkte projiziert werden. Dieser Ansatz wird spätestens dann suboptimal, wenn die Gruppen nicht

Abbildung 14: Klassifikation nach Fisher (1936) (I): Ω_1 blau, Ω_2 rot, eine mögliche Trennlinie T , eine Projektionsgerade Y



linear trennbar sind. Denn eigentlich braucht eine Trennebene nur durch Punkte definiert zu werden, an denen die Gruppen getrennt werden. Diese Punkte bzw. Vektoren sind die *support vectors*. In der nebenstehenden Abbildung sind die Punkte x_1 und x_2 Beispiele für support vectors. Sie liegen auf Rändern (margins), das sind Ebenen (hier: Geraden), die parallel zu einer "eigentlichen" Trennebene bzw. -geraden liegen. Die Trennung ist um so besser, je weiter die Margins von G entfernt liegen. Für einen gegebenen Datenpunkt $\mathbf{x}_i$ wird eine Funktion f gesucht derart, dass



$$y_i = f(\mathbf{x}_i) = \begin{cases} \geq +1, & \mathbf{x}_i \in \mathcal{C} \\ \leq -1, & \mathbf{x}_i \in \mathcal{C}^c \end{cases} \quad (7.1)$$

$$\mathbf{x} \mapsto f_{\mathbf{w},b}(\mathbf{x}) = \text{sgn}(\langle \mathbf{w}, \mathbf{x} \rangle + b) \quad (7.2)$$

gilt; sgn steht für *signum*, soll heißen für die Vorzeichen $+$ bzw. $-$; man kann auch die Faktoren $+1$ bzw. -1 damit bezeichnen. Die Gerade G und die zu ihr parallelen Geraden werden nur durch Punkte bestimmt, die auf den parallelen Geraden liegen, wie $\mathbf{x}_1$ und $\mathbf{x}_2$. Die Frage ist nun, wie diese Geraden bestimmt werden können. Eine Trennebene wird stets durch eine lineare Gleichung

$$w_1x_1 + w_2x_2 + \dots + w_px_p + b = 0$$

beschrieben. Eine Geradengleichung ergibt sich für $p = 2$. Es sei $\mathbf{w} = (w_1, \dots, w_p)'$; die Ebenen- bzw Geradengleichung für G ist dann durch

$$\mathbf{w}'\mathbf{x} + b = \langle \mathbf{w}, \mathbf{x} \rangle + b = 0 \quad (7.3)$$

gegeben, wobei b der Abstand der Ebene vom Nullpunkt des Koordinatensystems ist; die Schreibweise $\langle \mathbf{w}, \mathbf{x} \rangle$ für das Skalarprodukt ist in Texten zu SVMs üblich und wird deshalb hier übernommen, um den Übergang zur Originalliteratur zu erleichtern.

Der Vektor $\mathbf{w}$ ist orthogonal zur Orientierung der Ebene (Satz 2.2, Seite 42). Die Parallelen zu G werden durch die Forderung $\min_i \langle \mathbf{w}, \mathbf{x}_i \rangle + b = \pm 1$ definiert, d.h. durch

$$\min_i |\langle \mathbf{w}, \mathbf{x}_i \rangle + b| = 1 \quad (7.4)$$

definiert. Damit wird gesagt, dass der Punkt mit dem kleinsten Abstand zu G den Abstand $1/\|\mathbf{w}\|$ hat. Dass diese Aussage gilt, ist leicht zu sehen. Man betrachte zwei Punkte $\mathbf{x}_1$ und $\mathbf{x}_2$ mit $\langle \mathbf{w}, \mathbf{x}_1 \rangle + b = +1$ und $\langle \mathbf{w}, \mathbf{x}_2 \rangle + b = -1$, d.h. $|\langle \mathbf{w}, \mathbf{x} \rangle + b| = 1$. Subtrahiert man die zweite Gleichung von der ersten, so erhält man

$$\langle \mathbf{w}, (\mathbf{x}_1 - \mathbf{x}_2) \rangle = 2.$$

Normiert man den Vektor $\mathbf{w}$, d.h. geht man von $\mathbf{w}$ zu $\mathbf{w}/\|\mathbf{w}\|$ über, so erhält man

$$\langle \frac{\mathbf{w}}{\|\mathbf{w}\|}, \mathbf{x} \rangle = \frac{2}{\|\mathbf{w}\|}, \quad (7.5)$$

und dies bedeutet, dass der Abstand eines Margins zu G gerade gleich $1/\|\mathbf{w}\|$ ist.

Die Gerade G separiert optimal, wenn der Margin-Abstand maximal ist. Dies ist der Fall, wenn $1/\|\mathbf{w}\|$ maximal, d.h. wenn $\|\mathbf{w}\|$ minimal ist. $\mathbf{w}$ kann also durch Minimalisierung von $\|\mathbf{w}\|$ bestimmt werden, wobei allerdings Nebenbedingungen erfüllt sein müssen: es soll ja $|\langle \mathbf{w}, \mathbf{x} \rangle + b| \geq 1$ gelten. Diese Forderung kann man in der Form $y_i(\langle \mathbf{w}, \mathbf{x} \rangle + b) \geq 1$ schreiben: ist $\langle \mathbf{w}, \mathbf{x} \rangle + b > 0$, so ist auch $y_i > 0$ und ist $\langle \mathbf{w}, \mathbf{x} \rangle + b < 0$, so ist auch $y_i < 0$ und das Produkt der beiden Größen ist positiv. Damit besteht die Aufgabe der Bestimmung von $\mathbf{w}$ also darin, $\mathbf{w}$ gemäß

$$\min_{\mathbf{x} \in \mathcal{H}, b \in \mathbb{R}} \tau(\mathbf{w}) = \frac{1}{2} \|\mathbf{w}\|^2 \quad (7.6)$$

$$y_i(\langle \mathbf{w}, \mathbf{x} \rangle + b) \geq 1, \text{ für alle } i \quad (7.7)$$

zu bestimmen; der Faktor $1/2$ in (7.6) hat den Sinn, den Faktor 2 loszuwerden, der beim Differenzieren von $\tau(\mathbf{w})$ entsteht. Die Aufgabe wird gelöst, indem man die den Gleichungen (7.6) und (7.7) entsprechende Lagrange-Funktion

$$L(\mathbf{w}, b, \boldsymbol{\alpha}) = \frac{1}{2} \|\mathbf{w}\|^2 = \sum_{i=1}^m \alpha_i (y_i(\langle \mathbf{w}, \mathbf{x}_i \rangle + b) - 1) \quad (7.8)$$

betrachtet, wobei $\boldsymbol{\alpha}(\alpha_1, \dots, \alpha_m)'$ der Vektor der Lagrange-Multiplikatoren ist mit $\alpha_i \geq 0$. $L(\mathbf{w}, b, \boldsymbol{\alpha})$ muß bezüglich der α_i maximalisiert und bezüglich $\mathbf{w}$ minimiert werden. Dazu muß

$$\frac{\partial L}{\partial b} L(\mathbf{w}, b, \boldsymbol{\alpha}) = 0, \quad \frac{\partial L}{\partial \mathbf{w}} = 0 \quad (7.9)$$

gelten, woraus

$$\sum_{i=1}^m \alpha_i y_i = 0, \quad \mathbf{w} = \sum_{i=1}^m \alpha_i y_i \mathbf{x}_i \quad (7.10)$$

folgt (Schölkopf und Smola (2002), p. 197). Der gesuchte Vektor $\mathbf{w}$ ergibt sich demnach als Linearkombination der Vektoren $\mathbf{x}_i$, also aus den Daten der Trainingsstichprobe. Es läßt sich zeigen, dass nur diejenigen $\alpha_i \neq 0$ den Bedingungen von (7.6) und (7.7) genügen, für die

$$\alpha_i (y_i (\langle \mathbf{w}, \mathbf{x}_i \rangle + b) - 1) = 0 \quad (7.11)$$

gilt. Diejenigen $\mathbf{x}_i$ mit $\alpha_i > 0$ sind die *Support Vectors* (SVs)⁴³, denn (7.11) bedeutet ja gerade, dass diese $\mathbf{x}_i$ exakt auf dem Rand (Margin), also auf den Parallelen zu G liegen. Alle übrigen $\mathbf{x}_j$ sind irrelevant, da für sie $\alpha_i = 0$. Die Bedingungen (7.7) sind automatisch erfüllt.

Um das eigentliche Optimierungsproblem zu lösen, müssen die Bedingungen von (7.10) in die Lagrange-Funktion (7.8) eingesetzt werden. Man erhält

$$\max_{\boldsymbol{\alpha} \in \mathbb{R}^n} W(\boldsymbol{\alpha}) = \sum_{i=1}^m \alpha_i - \frac{1}{2} \sum_{i,j=1}^m \alpha_i \alpha_j \langle \mathbf{x}_i, \mathbf{x}_j \rangle \quad (7.12)$$

unter den Nebenbedingungen

$$\alpha_i \geq 0, \quad \sum_{i=1}^m \alpha_i y_i = 0. \quad (7.13)$$

Setzt man nun die Definition von $\mathbf{w}$ in die Entscheidungsfunktion (7.2) ein, so erhält man für f den Ausdruck

$$f(\mathbf{x}) = \text{sgn} \left(\sum_{i=1}^m \alpha_i y_i \langle \mathbf{x}, \mathbf{x}_i \rangle + b \right) \quad (7.14)$$

Damit wird deutlich, dass $f(\mathbf{x})$ u.a. durch die Skalarprodukte $\langle \mathbf{x}_i, \mathbf{x} \rangle$ definiert wird.

Dieser Sachverhalt ermöglicht den Übergang zu nichtlinearen Trennfunktionen, denn $\langle \mathbf{x}_i, \mathbf{x} \rangle$ ist der Spezialfall einer allgemeinen *Kernfunktion* $k(\mathbf{x}_i, \mathbf{x})$. Dazu wird das Skalar- oder innere Produkt $\langle \mathbf{x}_i, \mathbf{x} \rangle$ ersetzt:

$$\langle \mathbf{x}_i, \mathbf{x} \rangle \rightarrow \langle \phi(\mathbf{x}_i), \phi(\mathbf{x}) \rangle. \quad (7.15)$$

ϕ definiert den Merkmalsraum oder *feature space*, der im Allgemeinen von höherer Dimension als der ursprünglich gegebene ist und in dem die Trennung der

⁴³Man sollte vielleicht besser von Stützvektoren reden, um das unselige Denglisch zu vermeiden. Andererseits ist der englische Ausdruck zum Standardausdruck geworden, so dass er hier beibehalten wird

Gruppen durch lineare Funktionen möglich ist. Der Punkt hierbei ist, dass ϕ gar nicht explizit definiert und damit die Transformation in den höheren Raum nicht durchgeführt werden muß, da ϕ implizit durch die Wahl einer Kernfunktion gegeben ist:

$$k(\mathbf{x}_i, \mathbf{x}) = \langle \phi(\mathbf{x}_i), \phi(\mathbf{x}) \rangle. \quad (7.16)$$

Beispiele für Kernfunktionen sind

$$K_1(\mathbf{x}_i, \mathbf{x}_j) = \exp\left(-\frac{\|\mathbf{x}_i - \mathbf{x}_j\|}{2\sigma^2}\right) \quad (7.17)$$

$$K_2(\mathbf{x}_i, \mathbf{x}_j) = (\mathbf{x}_i, \mathbf{x}_j + 1)^d \quad (7.18)$$

$$K_3(\mathbf{x}_i, \mathbf{x}_j) = \tanh(\beta \mathbf{x}_i' \mathbf{x}_j + b) \quad (7.19)$$

Eine Anwendung und weitere Diskussion der SVM-Klassifikation findet man in Mortensen: Klassifikations- und Diskriminanzanalyse.

7.2 Kern-Regression

Geht es um Klassifikation, so wird – im dichotomen Fall – die abhängige Variable $y_i \in \{\pm 1\}$ gesetzt. Im Falle der Regressions nimmt y_i reelle Werte aus einem Intervall an. Der support-vector -Ansatz kann allerdings auf den Regressionsfall verallgemeinert werden. Es sei $f(\mathbf{x}) = \hat{y}$ die "vorhergesagte" Funktion, während die y -Werte die tatsächlich gemessenen Werte seien. Man definiert eine Verlustfunktion

$$c(\mathbf{x}, y, f(\mathbf{x})) = |y - f(\mathbf{x})|_\varepsilon = \max\{0, |y - f(\mathbf{x}) - \varepsilon|\}. \quad (7.20)$$

Es soll die Regressionsfunktion

$$f(\mathbf{x}) = \langle \mathbf{w}, \mathbf{x} \rangle + b = w_1 \mathbf{x}_1 + \dots + w_n \mathbf{x}_n + b \quad (7.21)$$

geschätzt werden. Dazu soll

$$\mathcal{L} = \frac{1}{2} \|\mathbf{w}\|^2 + C \sum_{i=1}^m |y_i - f(\mathbf{x}_i)|_\varepsilon \quad (7.22)$$

minimalisiert werden. Dazu werden "slack-Variablen" ξ_i und $\hat{\xi}_i$ eingeführt:

$$y_i - \langle \mathbf{w}, \mathbf{x} \rangle - b \leq \varepsilon + \xi_i, \quad (\langle \mathbf{w}, \mathbf{x} \rangle - y_i \leq \varepsilon + \hat{\xi}_i) \quad (7.23)$$

eingeführt und (7.22) wird modifiziert zu

$$\mathcal{L} = \min \left[\|\mathbf{x}\|^2 + C \sum_{i=1}^m (\xi_i + \hat{\xi}_i) \right] \quad (7.24)$$

mit $\xi_i \geq 0$, $\hat{\xi}_i \geq 0$. $\mathcal{L}$ ist eine *objective function*, d.h. eine Funktion, die maximiert bzw. minimalisiert werden soll. Die zu (7.24) duale objektive Funktion ist

$$\begin{aligned} W(\alpha, \alpha^*) &= \sum_{i=1}^m y_i(\alpha_i + \hat{\alpha}_i) - \varepsilon \sum_{i=1}^m (\alpha_i + \hat{\alpha}_i) \\ &\quad - \frac{1}{2} \sum_{i,j=1}^m (\alpha_i - \hat{\alpha}_i)(\alpha_j - \hat{\alpha}_j) K(\mathbf{x}_i, \mathbf{x}_j), \end{aligned} \quad (7.25)$$

und diese Funktion soll unter der Nebenbedingung

$$\sum_{i=1}^m \hat{\alpha}_i = \sum_{i=1}^m \alpha_i, \quad 0 \leq \alpha_i \leq C, \quad 0 \leq \hat{\alpha}_i \leq C \quad (7.26)$$

maximiert werden.

7.3 Kernel-PCA

Die PCA geht üblicherweise von einer Kovarianz- bzw. Korrelationsmatrix

$$C = (\mathbf{x}'_i \mathbf{x}_j) \quad (7.27)$$

aus. Das setzt eine lineare Beziehung zwischen den $\mathbf{x}_i, \mathbf{x}_j$ voraus. Es ist aber möglich, dass die Zusammenhänge durch Feature-Vektoren $\phi(\mathbf{x}_i), \phi(\mathbf{x}_j)$ definiert sind. Man kann dann C durch

$$C_\phi = \frac{1}{m} \sum_{i=1}^m \phi(\mathbf{x}_i) \phi(\mathbf{x}_i)' \quad (7.28)$$

ersetzen. Gesucht sind wieder die Lösungen $C_\phi \mathbf{v} = \lambda \mathbf{v}$. Der Eigenvektor $\mathbf{v}$ kann in eine Reihe

$$\mathbf{v} = \sum_{i=1}^m \alpha_i \phi(\mathbf{x}_i) \quad (7.29)$$

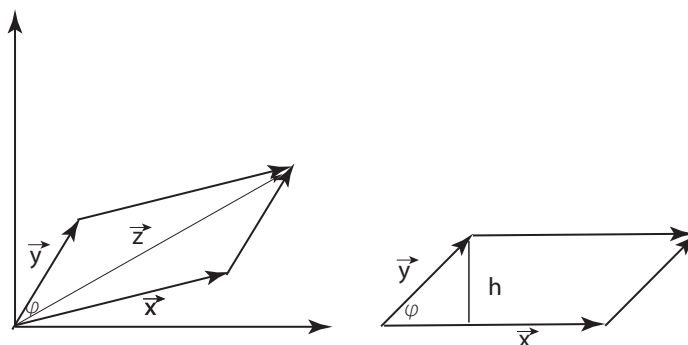
entwickelt werden, d.h. es müssen nun nur die α_i bestimmt werden. Man kann nun das Skalarprodukt zwischen $\phi(\mathbf{x})$ und dem n -ten normalisierten Eigenvektor berechnen:

$$\langle \mathbf{v}_n, \phi(\mathbf{x}) \rangle = \sum_{i=1}^m \alpha_i^n k(\mathbf{x}_i, \mathbf{x}) \quad (7.30)$$

$k(\mathbf{x}_i, \mathbf{x})$ eine Kernfunktion. Damit kann die explizite Berechnung der $\phi(\mathbf{x}_i)$ umgangen werden.

Ausführlichere Darstellungen werden in einem gesonderten Skript gegeben.

Abbildung 15: Zum vektoriellen Produkt: Parallelogramm und Fläche



8 Anhang

8.1 Das vektorielle Produkt

Es wird zunächst eine formale Definition dieses Produkts gegeben, die dann im Folgenden elaboriert wird.

Definition 8.1 Gegeben seien zwei 3-dimensionale Vektoren $\mathbf{x}$ und $\mathbf{y}$, die einen Winkel ϕ einschließen. Das vektorielle Produkt von $\mathbf{x}$ und $\mathbf{y}$ ist der durch

$$\mathbf{x} \times \mathbf{y} = \mathbf{z} = (\|\mathbf{x}\| \|\mathbf{y}\| \sin \phi) \mathbf{n} \quad (8.1)$$

definierte Vektor $\mathbf{z}$ der Länge $\|\mathbf{x}\| \|\mathbf{y}\| \sin \phi$, wobei $\|\mathbf{x}\|$ und $\|\mathbf{y}\|$ die Längen von $\mathbf{x}$ und $\mathbf{y}$ sind und der Vektor $\mathbf{n}$ ein zu $\mathbf{x}$ und $\mathbf{y}$ orthogonaler Vektor der Länge 1 ist.

Anmerkung: Das vektorielle Produkt wird hier nur für den üblichen 3-dimensionalen Fall diskutiert, Verallgemeinerungen sind möglich, gehen aber über den Rahmen dieses Textes hinaus. $\square$

Es wird zunächst gezeigt, dass der Faktor $\|\mathbf{x}\| \|\mathbf{y}\| \sin \phi$ von $\mathbf{n}$ gleich dem Flächeninhalt des durch $\mathbf{x}$ und $\mathbf{y}$ aufgespannten Parallelogramms ist. Anschließend werden die Komponenten des Vektors $\mathbf{z}$ hergeleitet.

Der Flächeninhalt eines Parallelogramms ist bekanntlich durch das Produkt $h\mathbf{x}$ gegeben (vergl. Abbildung 15). Der Sinus des Winkels ϕ ist durch den Quotienten Gegenkathete/Hypotenuse gegeben, also durch

$$\sin \phi = \frac{h}{\|\mathbf{y}\|}, \quad (8.2)$$

wobei $\|\mathbf{y}\|$ die Länge des Vektors $\mathbf{y}$ ist. Es folgt

$$h = \|\mathbf{y}\| \sin \phi. \quad (8.3)$$

Damit hat man für die Fläche des Parallelogramms den Ausdruck

$$F = \|\mathbf{x}\| \|\mathbf{y}\| \sin \phi. \quad (8.4)$$

Das vektorielle Produkt ist also durch $\mathbf{x} \times \mathbf{y} = \mathbf{z} = F \sin \phi \mathbf{n}$ gegeben und $F \sin \phi$ ist gleich der Länge $\|\mathbf{z}\|$ von $\mathbf{z}$, die mithin vom Winkel ϕ abhängt, den die Vektoren $\mathbf{x}$ und $\mathbf{y}$ einschließen. Für $\phi = 0$ ist $\sin \phi = 0$, – haben also $\mathbf{x}$ und $\mathbf{y}$ dieselbe Orientierung, so ist $F \sin \phi = 0$ und $\mathbf{z}$, d.h. das vektorielle Produkt, hat den Wert 0. Für gegebene Längen $\|\mathbf{x}\|$ und $\|\mathbf{y}\|$ wird $F \sin \phi$ maximal, wenn $\phi = \pi/2$, denn dann wird $\sin \phi = 1$. Für $\phi = 3\pi/2$ ist $\sin \phi = -1$, $\|\mathbf{z}\|$ nimmt wieder einen maximalen Wert an, allerdings zeigt $\mathbf{z}$ nun in die entgegengesetzte Richtung (wie für alle Werte von ϕ zwischen π und 2π).

Die Komponenten von $\mathbf{z}$: Bis jetzt ist noch nicht erklärt worden, wie der Vektor $\mathbf{z}$ tatsächlich bestimmt werden kann. Diese Bestimmung ergibt sich, wenn man den Winkel ϕ aus dem Ausdruck für $\mathbf{x} \times \mathbf{y}$ eliminiert, was schon deswegen eine nützliche Übung ist, weil er oft gar nicht gegeben ist. Nun ist $\sin^2 \phi + \cos^2 \phi = 1$, so dass $\sin \phi = \sqrt{1 - \cos^2 \phi}$ folgt, und wegen $\cos \phi = \mathbf{x}'\mathbf{y}/(\|\mathbf{x}\| \|\mathbf{y}\|)$ (vergl. (2.31), Seite 23) kann $\sin \phi$ durch einen nur von den Vektoren $\mathbf{x}$ und $\mathbf{y}$ zusammengesetzten Ausdruck ersetzt werden: man hat

$$\begin{aligned} F &= \|\mathbf{x}\| \|\mathbf{y}\| \sqrt{1 - \cos^2 \phi} \\ &= \|\mathbf{x}\| \|\mathbf{y}\| \sqrt{1 - \frac{(\mathbf{x}'\mathbf{y})^2}{\|\mathbf{x}\|^2 \|\mathbf{y}\|^2}} = \|\mathbf{x}\| \|\mathbf{y}\| \sqrt{\frac{\|\mathbf{x}\|^2 \|\mathbf{y}\|^2 - (\mathbf{x}'\mathbf{y})^2}{\|\mathbf{x}\|^2 \|\mathbf{y}\|^2}} \\ &= \sqrt{\|\mathbf{x}\|^2 \|\mathbf{y}\|^2 - (\mathbf{x}'\mathbf{y})^2} \\ &= \sqrt{(x_1^2 + x_2^2 + x_3^2)(y_1^2 + y_2^2 + y_3^2) - (x_1 y_1 + x_2 y_2 + x_3 y_3)^2} \quad (8.5) \end{aligned}$$

Mit (8.5) hat man einen Ausdruck für die Fläche F , in der der Winkel ϕ nicht mehr vorkommt. Es zeigt sich nun, dass der Ausdruck

$$A = (x_1^2 + x_2^2 + x_3^2)(y_1^2 + y_2^2 + y_3^2) - (x_1 y_1 + x_2 y_2 + x_3 y_3)^2$$

unter der Wurzel in (8.5) sich so umformen läßt, dass sich eine Charakterisierung des Vektors $\mathbf{z}$ ergibt. Multipliziert man A aus, so ergibt sich

$$\begin{aligned} A &= (x_1 y_1)^2 + (x_1 y_2)^2 + (x_1 y_3)^2 + (x_2 y_1)^2 + (x_2 y_2)^2 + (x_2 y_3)^2 \\ &\quad + (x_3 y_1)^2 + (x_3 y_2)^2 + (x_3 y_3)^2 \\ &\quad - (x_1 y_1)^2 - (x_2 y_2)^2 - (x_3 y_3)^2 - 2(x_1 y_1 x_2 y_2) - 2(x_1 y_1 x_3 y_3) - 2(x_2 y_2 x_3 y_3) \\ &= (x_1 y_2)^2 + (x_2 y_1)^2 + (x_1 y_3)^2 + (x_3 y_1)^2 + (x_2 y_3)^2 + (x_3 y_2)^2 \\ &\quad - 2(x_1 y_1 x_2 y_2) - 2(x_1 y_1 x_3 y_3) - 2(x_2 y_2 x_3 y_3) \end{aligned}$$

Die Inspektion dieser Summe legt nahe, die Terme $(x_1 y_2)^2$ und $(x_2 y_1)^2$, $(x_1 y_3)^2$, $(x_3 y_1)^2$ und $(x_2 y_3)^2$, $(x_3 y_2)^2$ zusammenfassen:

$$\begin{aligned} (x_1 y_2)^2 + (x_2 y_1)^2 &= (x_1 y_2 - x_2 y_1)^2 + 2x_1 y_2 x_2 y_1 \\ (x_1 y_3)^2 + (x_3 y_1)^2 &= (x_1 y_3 - x_3 y_1)^2 + 2x_1 y_3 x_3 y_1 \\ (x_2 y_3)^2 + (x_3 y_2)^2 &= (x_2 y_3 - x_3 y_2)^2 + 2x_2 y_3 x_3 y_2 \end{aligned}$$

Die Terme $2x_1y_2x_2y_1$ etc heben sich dann im Ausdruck für A heraus und man erhält

$$A = (x_1y_2 - x_2y_1)^2 + (x_2y_3 - x_3y_2)^2 + (x_1y_3 - x_3y_1)^2.$$

Demnach kann $F = \sqrt{A}$ als Länge eines Vektors $\mathbf{u}_0$ mit den Komponenten $a = x_1y_2 - x_2y_1$, $b = x_2y_3 - x_3y_2$ und $c = x_1y_3 - x_3y_1$ aufgefasst werden. Aber A ist durch die Quadratsumme dieser Komponenten definiert, so dass nicht nur $A = \|\mathbf{u}_0\|$, sondern ebenso $A = \|\mathbf{u}_j\|$ gilt, wobei $\mathbf{u}_j$ durch eine der Permutationen der Komponenten a , b und c definiert ist. Alle diese Vektoren haben die gleiche Länge, aber sind durch verschiedene Orientierungen charakterisiert. Ebenso können die Vorzeichen der a , b und c variiert werden, dabei bleiben die Längen ebenfalls invariant. F soll aber die Länge eines Vektors $\mathbf{z}$ sein, der orthogonal zu $\mathbf{x}$ und $\mathbf{y}$ ist, so dass aus der Menge der möglichen Vektoren derjenige ausgewählt werden muß, der diese Bedingung der Orthogonalität erfüllt. Man wählt dazu diejenige Kombination von Vorzeichen und Anordnung der Komponenten a , b und c , dass ein Vektor entsteht, der auf den Vektoren $\mathbf{x}$ und $\mathbf{y}$ senkrecht steht. Setzt man also $\mathbf{z} = (b_{j_1}, b_{j_2}, b_{j_3})'$, wobei $(b_{j_1}, b_{j_2}, b_{j_3})$ eine bestimmte Anordnung der a, b, c mit bestimmten Vorzeichen ist, so soll

$$(b_{j_1}, b_{j_2}, b_{j_3}))\mathbf{x}' = b_{j_1}x_1 + b_{j_2}x_2 + b_{j_3}x_3 = 0$$

gelten. Man findet dann, dass die Wahl

$$b_{j_1} = x_2y_3 - x_3y_2, \quad b_{j_2} = x_3y_1 - x_1y_3, \quad b_{j_3} = x_1y_2 - x_2y_1$$

für die Komponenten von $\mathbf{z}$, also

$$\mathbf{z} = \begin{pmatrix} x_2y_3 - x_3y_2 \\ x_3y_1 - x_1y_3 \\ x_1y_2 - x_2y_1 \end{pmatrix} \quad (8.6)$$

gerade die gewünschte Orthogonalitätsforderung erfüllt. Um das zu sehen, müssen nur die Skalarprodukte $\mathbf{z}'\mathbf{x}$ und $\mathbf{z}'\mathbf{y}$ betrachtet werden. Es ist

$$\mathbf{z}'\mathbf{x} = (x_2y_3 - x_3y_2)x_1 + (x_3y_1 - x_1y_3)x_2 + (x_1y_2 - x_2y_1)x_3,$$

und durch Ausmultiplizieren findet man

$$x_1x_2y_3 - x_1x_3y_2 + x_2x_3y_1 - x_1x_2y_3 + x_1x_3y_2 - x_2x_3y_1 = 0$$

Die Orthogonalität von $\mathbf{z}$ und $\mathbf{y}$ folgt analog. Der so charakterisierte Vektor entspricht einem Rechtssystem im oben definierten Sinn.

Mithin kann das vektorielle Produkt in der Form

$$\mathbf{x} \times \mathbf{y} = \|\mathbf{x}\|\|\mathbf{y}\| \sin \phi \mathbf{n} = \mathbf{z} = \begin{pmatrix} x_2y_3 - x_3y_2 \\ x_3y_1 - x_1y_3 \\ x_1y_2 - x_2y_1 \end{pmatrix}. \quad (8.7)$$

geschrieben werden.

Die Orientierung von $\mathbf{z}$ kann mit einer Drehrichtung verbunden werden. Dazu werde angenommen, dass man $\mathbf{x}$ mit dem Daumen der *rechten* Hand assoziiert und den Vektor $\mathbf{y}$ mit dem Zeigefinger ebenfalls der rechten Hand. Nun führe man mit der Hand eine Rechtsdrehung aus, wie man sie beim Eindrehen einer Schraube mit dem üblichen Rechtsgewinde durchführt, – es ist, als wolle man $\mathbf{x}$ um den Winkel ϕ in die Richtung von $\mathbf{y}$ drehen. $\mathbf{z}$ zeigt dann in Richtung der Schraube. Man spricht von einem *Rechtssystem*. Das *Linkssystem* ist analog definiert: man rotiert den Vektor $\mathbf{y}$ um den Winkel $-\phi$ in die Richtung von $\mathbf{x}$. Die meisten Anwendungen findet man in der Physik (z.B. in der Mechanik und in der Elektrodynamik).

Die Spezifikation von $\mathbf{z}$ in (8.6) ist gleichwohl nicht eindeutig, denn auch $-\mathbf{z}$ steht senkrecht auf $\mathbf{x}$ und $\mathbf{y}$ bzw. auf der durch diese Vektoren aufgespannten Ebene, zeigt aber in die entgegengesetzte Richtung von $\mathbf{z}$ und definiert ein Linkssystem.

Eine elegante Möglichkeit, das vektorielle Produkt zu definieren und seine Eigenschaften herzuleiten, besteht darin, es symbolisch als Determinante zu definieren und diese nach der ersten Zeile zu entwickeln:

$$\begin{aligned}\mathbf{x} \times \mathbf{y} = \mathbf{z} &= \begin{vmatrix} \mathbf{e}_1 & \mathbf{e}_2 & \mathbf{e}_3 \\ x_1 & x_2 & x_3 \\ y_1 & y_2 & y_3 \end{vmatrix} \\ &= (x_2y_3 - x_3y_2)\mathbf{e}_1 - (x_1y_3 - x_3y_1)\mathbf{e}_2 + (x_1y_2 - x_2y_1)\mathbf{e}_3, \quad (8.8)\end{aligned}$$

wobei die Vorzeichen der Summanden nach der Regel $(-1)^{1+j}$ bestimmt wurden. Der Ansatz macht davon Gebrauch, dass ein beliebiger Vektor $\mathbf{v} = (v_1, v_2, v_3)'$ in der Form $\mathbf{v} = v_1\mathbf{e}_1 + v_2\mathbf{e}_2 + v_3\mathbf{e}_3$ dargestellt werden kann. Tatsächlich entsprechen die Koeffizienten der $\mathbf{e}_j$ den Komponenten des Vektors $\mathbf{z}$ in (8.7), wenn man $-(x_1y_3 - x_3y_1) = x_3y_1 - x_1y_3$ berücksichtigt.

Es gilt der

Satz 8.1 *Es gelten die Aussagen*

1. $\mathbf{x} \times \mathbf{y} = \mathbf{0}$ genau dann, wenn entweder (i) $\mathbf{x} = \vec{0}$ oder $\mathbf{y} = \vec{0}$, oder (ii) $\mathbf{x}$ und $\mathbf{y}$ parallel sind,
2. $\mathbf{x} \times \mathbf{y} = -\mathbf{y} \times \mathbf{x}$,
3. $\mathbf{x} \times (\mathbf{y} + \mathbf{z}) = \mathbf{x} \times \mathbf{y} + \mathbf{x} \times \mathbf{z}$. (*Distributivgesetz*)

Beweis: Die Aussage 1. (i) ist klar: ist entweder $\mathbf{x} = \vec{0}$ oder $\mathbf{y} = \vec{0}$, so sind die Koeffizienten von $\mathbf{z}$ alle gleich Null. Haben $\mathbf{x}$ und $\mathbf{y}$ dieselbe Orientierung, so gilt $\mathbf{y} = \lambda\mathbf{x}$ für $\lambda \neq 0$ und man hat $x_2y_3 - x_3y_2 = \lambda x_2x_3 - \lambda x_3x_2 = 0$ etc.

Die Aussage 2. impliziert, dass sich die Vorzeichen der Koeffizienten von $\mathbf{z}$ ändern bzw. die Koeffizienten alle mit dem Faktor -1 multipliziert werden. Dieser

Fall resultiert, wenn man $\mathbf{y} \times \mathbf{x}$ berechnet:

$$\begin{aligned}\mathbf{y} \times \mathbf{x} = \mathbf{z} &= \begin{vmatrix} \mathbf{e}_1 & \mathbf{e}_2 & \mathbf{e}_3 \\ y_1 & y_2 & y_3 \\ x_1 & x_2 & x_3 \end{vmatrix} \\ &= (y_2x_3 - y_3x_2)\mathbf{e}_1 - (y_1x_3 - y_3x_1)\mathbf{e}_2 + (y_1x_2 - y_2x_1)\mathbf{e}_3 \\ &= -(x_2y_3 - x_3y_2)\mathbf{e}_1 + (x_1y_3 - x_3y_1)\mathbf{e}_2 - (x_1y_2 - x_2y_1)\mathbf{e}_3 \\ &= -\mathbf{x} \times \mathbf{y}\end{aligned}$$

Der Vergleich mit dem Ausdruck für $\mathbf{x} \times \mathbf{y}$ vergleicht bestätigt die Aussage.

Die Aussage 3. bestätigt man, indem man $\mathbf{u} = \mathbf{y} + \mathbf{z}$ setzt und das vektorielle Produkt $\mathbf{x} \times \mathbf{u}$ bestimmt. $\square$

8.2 Ungleichungen

8.2.1 Eine allgemeine Ungleichung

Es sei $f(x)$ eine Funktion, für die für alle x aus einem Bereich $I = (x_1, x_2)$ die Ableitung $df/dx = f'(x)$ existiert. Wird auf f' auf I kleiner, je größer x wird, so heißt die Funktion auf I *konkav*, wächst dagegen $f'(x)$ mit größer werdendem $x \in I$, so heißt die Funktion auf I *konvex*. Der Logarithmus $f(x) = \log x$ ist auf dem gesamten Intervall $0 < x < \infty$ konkav: $\log x$ wächst monoton mit x , aber mit größer werdendem x immer langsamer, denn $f'(x) = 1/x$, $x > 0$, und offenbar $f'(x) \rightarrow 0$ für $x \rightarrow \infty$.

Es sei f eine konkave Funktion (über einem Intervall I). Es sei weiter $b > a$, so dass $f(b) > f(a)$. Weiter sei $x \in (a, b)$. Dann ist $f(a) < f(x) < f(b)$, und $f(x)$ liegt stets über der Geraden, die die Punkte $f(a)$ und $f(b)$ miteinander verbindet. Daraus folgt, dass

$$\frac{f(x) - f(a)}{x - a} \geq \frac{f(b) - f(a)}{b - a},$$

woraus wiederum

$$f(x) \geq f(a) + (f(b) - f(a)) \frac{x - a}{b - a} = f(a) + \lambda(f(b) - f(a)), \quad \lambda \in (0, 1)$$

so dass

$$f(x) \geq (1 - \lambda)f(a) + \lambda f(b), \quad \forall x \in (a, b). \quad (8.9)$$

Insbesondere kann

$$x = x(\lambda) = (1 - \lambda)a + \lambda b = a + \lambda(b - a)$$

gewählt werden. X ist eine lineare Funktion von λ mit der Steigung $dx/d\lambda = b - a$, und der additiven Konstante a , so dass $a \leq x \leq b$; für $\lambda = 0$ ist $x = a$, und für $\lambda = 1$ ist $x = b$. Da (8.9) allgemein gilt hat man dann

$$f(x) = f((1 - \lambda)a + \lambda b) \geq (1 - \lambda)f(a) + \lambda f(b). \quad (8.10)$$

Es sei nun insbesondere $f(x) = \log x$. Dann folgt

$$\log((1-\lambda)a + \lambda b) \geq (1-\lambda)\log a + \lambda\log b = \lambda_1 \log a + \lambda_2 \log b = \log(a^{\lambda_1} b^{\lambda_2}) \quad (8.11)$$

wobei $\lambda_1 = 1 - \lambda$, $\lambda_2 = \lambda$, so dass $\lambda_1 + \lambda_2 = 1$. Man kann diese Ungleichung sofort verallgemeinern zu

$$\log(\lambda_1 a_1 + \lambda_2 a_2 + \cdots + \lambda_n a_n) \geq \log(a_1^{\lambda_1} a_2^{\lambda_2} \cdots a_n^{\lambda_n}), \quad \sum_{i=1}^n \lambda_i = 1, \quad (8.12)$$

mit $a_i > 0$ für $i = 1, \dots, n$, und da der Logarithmus eine monotone Funktion ist, folgt

$$\lambda_1 a_1 + \lambda_2 a_2 + \cdots + \lambda_n a_n \geq a_1^{\lambda_1} a_2^{\lambda_2} \cdots a_n^{\lambda_n}. \quad (8.13)$$

Setzt man insbesondere $\lambda_1 = \cdots = \lambda_n = 1/n$, so erhält man die Ungleichung

$$\frac{1}{n} \sum_{i=1}^n a_i \geq \prod_{i=1}^n a_i^{1/n} = \left(\prod_{i=1}^n a_i \right)^{1/n}. \quad (8.14)$$

Links steht das arithmetische Mittel der a_i , rechts das geometrische Mittel der a_i , d.h. man hat hier die bekannte Ungleichung zwischen arithmetischem und geometrischem Mittel.

8.2.2 Die Höldersche Ungleichung

Die Ungleichung (8.13) führt auf eine weitere, wichtige Ungleichung:

Satz 8.2 (Höldersche Ungleichung) *Es gelte $p > 0$ und $1/p + 1/q = 1$. Dann gilt*

$$\sum_{i=1}^n |a_i b_i| \leq \left(\sum_{i=1}^n |a_i|^p \right)^{1/p} \left(\sum_{i=1}^n |b_i|^q \right)^{1/q}. \quad (8.15)$$

Beweis: A und B seien die Summen

$$A = \sum_{i=1}^n |a_i|^p, \quad B = \sum_{i=1}^n |b_i|^q. \quad (8.16)$$

Es werde $p = 1/\lambda_1 > 1$ und $\lambda_2 = 1/q$ gesetzt, so dass $\lambda_1 + \lambda_2 = 1$. Die zwei Terme $|a_i|/A$ und $|b_i|/B$ sind sicher positiv, ersetzt man a_1 in (8.13) durch $|a_i|/A$ und a_2 durch $|b_i|/B$, so liefert (8.13) die Ungleichung

$$\left(\frac{|a_i|^{1/\lambda_1}}{A} \right)^{\lambda_1} \left(\frac{|b_i|^{1/\lambda_2}}{B} \right)^{\lambda_2} \leq \lambda_1 \frac{|a_i|^{1/\lambda_1}}{A} + \lambda_2 \frac{|b_i|^{1/\lambda_2}}{B}, \quad 1 \leq i \leq n$$

Summiert man beide Seiten über i , so erhält man

$$\sum_{i=1}^n \left(\frac{|a_i|^{1/\lambda_1}}{A} \right)^{\lambda_1} \left(\frac{|b_i|^{1/\lambda_2}}{B} \right)^{\lambda_2} \leq \lambda_1 \frac{A}{A} + \lambda_2 \frac{B}{B} = \lambda_1 + \lambda_2 = 1$$

und man erhält

$$\sum_{i=1}^n \frac{|a_i|}{A^{\lambda_1}} \frac{|b_i|}{B^{\lambda_2}} \leq 1,$$

woraus sofort

$$\sum_{i=1}^n |a_i| |b_i| \leq A^{\lambda_1} B^{\lambda_2}$$

folgt und wegen der Definition von A, B in (8.16) hat man die Behauptung (8.15). $\square$

8.3 Der allgemeine Begriff des Vektorraums

Der Begriff des Vektorraums ist sehr allgemein und setzt den Begriff des Körpers voraus, der wiederum den Begriff der Gruppe voraussetzt.

Definition 8.2 Eine Gruppe ist eine Menge $\mathcal{G}$, für deren Elemente eine Verknüpfung $\otimes$ definiert ist. Dabei gilt:

1. Für $a, b \in \mathcal{G}$ ist $a \otimes b \in \mathcal{G}$,
2. $\otimes$ ist assoziativ, d.h. für $a, b, c \in \mathcal{G}$ gilt $(a \otimes b) \otimes c = a \otimes (b \otimes c)$,
3. Es existiert ein neutrales Element $e \in \mathcal{G}$ derart, dass für alle $a \in \mathcal{G}$ die Beziehung $a \otimes e = e \otimes a = a$ gilt.
4. Für jedes Element $a \in \mathcal{G}$ existiert ein inverses Element $a^{-1} \in \mathcal{G}$, so dass $a \otimes a^{-1} = a^{-1} \otimes a = e$. Die Gruppe heißt abelsch⁴⁴, wenn $\otimes$ kommutativ ist, d.h. wenn $a \otimes b = b \otimes a$ für $a, b \in \mathcal{G}$ ist.

Es sei $\mathbb{R}$ die Menge der reellen Zahlen. Man kann $\otimes = +$ (Addition), oder $\otimes = \cdot$ die Multiplikation oder $\otimes = \div$ die Division setzen; $\mathbb{R}$ erweist sich dann in Bezug auf diese Operationen als Gruppe. Weiter sei M eine Menge von Objekten im 3-dimensionalen Raum, und $\otimes$ sei eine Drehung um einen bestimmten Winkel θ . Die Menge der Drehungen eines Objekts aus M ist eine Gruppe.

Definition 8.3 Es sei K eine Menge mit den zwei Verknüpfungen:

$$+ : K \times K \rightarrow K, (a, b) \mapsto a + b$$

$$\cdot : K \times K \rightarrow K, (a, b) \mapsto a \cdot b, a, b \in K.$$

heißt Körper, wenn die folgenden drei Bedingungen erfüllt sind:

K1: K zusammen mit der Addition $+$ ist eine abelsche Gruppe; das neutrale Element wird mit 0 bezeichnet, das zu $a \in K$ inverse Element mit $-a$.

⁴⁴Nach dem norwegischen Mathematiker Niels Henrik Abel (1802 – 1829)

K2: Es sei $K^* = K \setminus \{0\}$, d.h. K^* sei K ohne das Nullelement. Für $a, b \in K^*$ ist auch $a \cdot b \in K^*$, und K^* ist eine abelsche Gruppe mit dem neutralen Element 1, das zu a inverse Element a^{-1} wird auch mit $1/a$ bezeichnet. Schreibweise: $b/a = a^{-1}b = ba^{-1}$.

K3: Für $a, b, c \in K$ gilt

$$a \cdot (b + c) = a \cdot b + a \cdot c,$$

d.h. es gilt das Distributivgesetz.

Offenbar bilden die Menge $\mathbb{R}$ der reellen Zahlen und die Menge $\mathbb{C}$ der komplexen Zahlen jeweils einen Körper. Aber auch Mengen von Funktionen über einem Intervall $[a, b]$ oder von Matrizen können einen Körper bilden. Der Sinn der Einführung des Gruppen- und des Körperbegriffs ist, dass Vektorräume für beliebige Mengen von Objekten definiert werden können, sofern für diese Objekte Verknüpfungen definiert sind, die den Gruppen- bzw. den Körperbedingungen genügen. In der multivariaten Statistik beschränkt man sich aber auf Vektoren, deren Komponenten reelle oder komplexe Zahlen sind und auf bestimmte Typen von Funktionen.

Es folgt nun die allgemeine Definition des Begriffs des Vektorraums:

Definition 8.4 Eine nicht leere Menge V heißt Vektorraum über einem Körper $\mathbb{K}$, wenn alle Linearkombinationen von Elementen von V wieder Elemente von V sind. Insbesondere ist V ein Vektorraum, wenn die Bedingungen

- (i) Ist $\mathbf{v} \in V$, so auch $-\mathbf{v} \in V$
- (ii) $\lambda, \mu \in \mathbb{R}$, $\mathbf{v}, \mathbf{w} \in V \Rightarrow (\lambda + \mu)\mathbf{v} = \lambda\mathbf{v} + \mu\mathbf{v} \in V$, und $\lambda(\mathbf{v} + \mathbf{w}) = \lambda\mathbf{v} + \lambda\mathbf{w} \in V$
- (iii) $\lambda(\mu\mathbf{v}) = (\lambda\mu)\mathbf{v} \in V$ für $\mathbf{v} \in V$.

erfüllt sind. Die Elemente eines Vektorraums heißen Vektoren.

Anmerkungen:

- Vektorbegriff:** Die Definition des Vektorraums ist ohne Spezifikation der Elemente formuliert worden, und dass dann diese Elemente "Vektoren" genannt werden, verweist darauf, dass nicht nur die bisher betrachteten Vektoren $\mathbf{v} = (v_1, \dots, v_n)'$ gemeint sind, – diese sind eher Spezialfälle. Alle mathematischen Objekte sind Vektoren, sofern die Verknüpfungsregeln (i) bis (iii) auf sie anwendbar sind. So sind bilden z.B. bestimmte Mengen von Matrizen einen Vektorraum, oder Mengen von Polynomen, allgemein von stetigen Funktionen, etc.
- n -dimensionaler Vektorraum:** Sind die Elemente von V n -dimensionale Vektoren $\mathbf{x} = (x_1, \dots, x_n)'$, so schreibt man für den entsprechenden Vektorraum auch V_n , sind die Komponenten überdies reelle Zahlen, so ist auch

die Schreibweise $\mathbb{R}^n$ gebräuchlich, oder $\mathbb{C}^n$, wenn die Komponenten komplexe Zahlen sind. Eine etwas anders lautende Definition des n -dimensionalen Vektorraums wird in Definition 2.16, Seite 47, gegeben. $\square$

Mit dem Begriff des Vektorraums wird die Idee der Abgeschlossenheit einer Menge in Bezug auf die Verknüpfungsoperationen charakterisiert. Beispiele für einen Vektorraum sind

1. Die Menge $\mathbb{R}^n = \mathbb{R} \times \cdots \times \mathbb{R}$ aller n -tupel $x_1, \dots, x_n$ reeller Zahlen über dem Körper $\mathbb{R}$, also die Menge der hier betrachteten n -dimensionalen Vektoren mit reellen Komponenten.
2. Die Menge $C^k[a, b]$ aller auf dem abgeschlossenen Intervall⁴⁵ $[a, b]$ k -mal stetig differenzierbarer Funktionen. Sind $f \in C^k$, $g \in C^k$, so gilt $(f+g)(t) = f(t) + g(t)$ und $(f \cdot g)(t) = f(t)g(t)$ für alle $t \in [a, b] \subset \mathbb{R}$.
3. Die Menge der auf dem Intervall $[0, \infty)$ stetigen Funktionen.

Beispiel 8.1 Ein Spezialfall der stetigen Funktionen sind Polynome. Es sei $M = \{1, t, t^2, t^3, \dots\} \subset V$ und weitere Elemente ("Vektoren") in V sind eben die Funktionen

$$P(t) = a_0 + a_1 t + a_2 t^2 + \cdots + a_n t^n + \cdots.$$

M ist linear unabhängig, wenn M unendlich viele Elemente enthält.

Auf den ersten Blick scheinen Polynome, oder allgemein stetige Funktionen, nicht viel mit den Vektoren $\mathbf{x} = (x_1, \dots, x_n)'$ gemein zu haben, die bisher betrachtet wurden. Andererseits muß man sich klar machen, dass mit $P(t)$ ja nicht ein einzelner Wert (berechnet für einen bestimmten t -Wert) gemeint ist, sondern eben die Funktion P auf einem Definitionsbereich für t . So sei einmal angenommen, man habe P für die Werte $t_1, t_2, \dots, t_n$ berechnet; man erhält dann Werte $x_1 = P(t_1), x_2 = P(t_2), \dots, x_n = P(t_n)$. Der endlichdimensionale Vektor $(x_1, \dots, x_n)'$ liefert dann ein erstes, unvollständiges Bild von $P(t)$ für $t \in I$, $I \subseteq \mathbb{R}$ ein Intervall. Würde man $P(t)$ für *alle* $t \in I$ berechnen, bekäme man einen Vektor mit überabzählbar⁴⁶ vielen Komponenten, $\mathbf{x} = P$ hat dann ein Kontinuum von

⁴⁵Ein Intervall heißt *abgeschlossen*, wenn die Endpunkte a und b mit zum Intervall gehören, es wird dann $[a, b]$ geschrieben; gehören sie nicht zum Intervall, so heißt das Intervall *offen* und es wird (a, b) geschrieben, und gehört nur einer der Endpunkte zum Intervall, so heißt es *halb offen*, was durch $[a, b)$ oder $(a, b]$ signalisiert wird.

⁴⁶Eine Menge M heißt *abzählbar*, wenn man ihre Elemente durchnummerieren kann. Das ist auf jeden Fall möglich, wenn M endlich viele Elemente enthält. Sie darf auch unendlich viele Elemente enthalten: so lange man jedem Element $x \in M$ eine natürliche Zahl $i \in \mathbb{N}$ zuordnen kann, heißt sie abzählbar: $M = \{x_1, x_2, x_3, \dots, x_i, \dots | i \in \mathbb{N}\}$, M hat dann die "Mächtigkeit" der Menge $\mathbb{N}$ der natürlichen Zahlen. Es läßt sich zeigen, dass auch die Menge $\mathbb{Q}$ der rationalen (von lat. ratio = Quotient) Zahlen, die in der Form p/q mit $p, q \in \mathbb{N}$, darstellbar sind, noch die Mächtigkeit von $\mathbb{N}$ hat, – die rationalen Zahlen lassen sich also durchnummerieren. Nimmt man allerdings noch die irrationalen Zahlen – Dezimalzahlen, die eben nicht in der Form p/q mit $p, q \in \mathbb{N}$ darstellbar sind, wie etwa $\sqrt{2}$, π , e . etc – so lässt sich die Elemente nicht mehr durchnummerieren, man spricht dann von einer *überabzählbaren* Menge.

Komponenten. In Abschnitt 2.4.4 wird noch einmal auf Polynome als Vektoren eingegangen. $\square$

Im Rahmen der gewöhnlichen multivariaten Analyse ist so gut wie immer die Menge $\mathbb{R}$ der reellen Zahlen der betrachtete Körper, und im Zusammenhang mit Analysen dynamischer Vorgänge noch der allgemeinere Körper $\mathbb{C}$ der komplexen Zahlen. Analysen, bei denen auf den Körper $C^k[a, b]$ Bezug genommen wird, werden später in einem gesonderten Kapitel behandelt. Es genügt für das Folgende, einfach von einem Vektorraum zu sprechen, ohne zu erwähnen, dass es sich dabei um einen Vektorraum über einem Körper $\mathbb{K}$ handeln soll, wobei $\mathbb{K} = \mathbb{R}$ oder $\mathbb{K} = \mathbb{C}$ ist.

Nicht jede Menge von Vektoren ist ein Vektorraum. So sei

$$V = \{\mathbf{x} \mid \|\mathbf{x}\| = 1\}$$

eine Menge von Vektoren der Länge 1; legt man die Anfangspunkte dieser Vektoren in den Ursprung des Koordinatensystems, so liegen die Endpunkte der Elemente von V auf einem Kreis. V ist kein Vektorraum, denn für $\lambda, \mu \in \mathbb{R}$ beliebig ist

$$\mathbf{x} = \lambda \mathbf{x}_1 + \mu \mathbf{x}_2, \quad \mathbf{x}_1, \mathbf{x}_2 \in V$$

ein Vektor, dessen Länge $\lambda + \mu$ nicht notwendig gleich 1 ist, d.h. $\mathbf{x}$ ist nicht notwendig ein Element aus V .

8.4 Einfache lineare Regression

Als Illustration der bisher eingeführten Begriffe kann die einfache lineare Regression und die Schätzung des Regressionsparameters betrachtet werden. Es gelte

$$y_i = bx_i + e_i, \quad i = 1, \dots, n \quad (8.17)$$

mit $x_i = X_i - \bar{x}$, $y_i = Y_i - \bar{y}$, so dass $\sum_i x_i = \sum_i y_i = 0$. Der Parameter b ist unbekannt und soll geschätzt werden.

Statt (8.17) kann die Gleichung vektoriell geschrieben werden:

$$\mathbf{y} = b\mathbf{x} + \mathbf{e}, \quad (8.18)$$

mit $\mathbf{y} = (y_1, \dots, y_n)'$, $\mathbf{x} = (x_1, \dots, x_n)'$ und $\mathbf{e} = (e_1, \dots, e_n)'$. Eine direkte Möglichkeit, b zu schätzen, besteht darin, die Gleichung in eine Gleichung zwischen Skalaren zu verwandeln. Dazu werde die Gleichung von links mit $\mathbf{x}'$ multipliziert:

$$\mathbf{x}'\mathbf{y} = b\mathbf{x}'\mathbf{x} + \mathbf{x}'\mathbf{e},$$

woraus sofort

$$\frac{\mathbf{x}'\mathbf{y}}{\mathbf{x}'\mathbf{x}} - \frac{\mathbf{x}'\mathbf{e}}{\mathbf{x}'\mathbf{x}} = \tilde{b} \quad (8.19)$$

folgt. Für den ersten Term auf der linken Seite findet man

$$\frac{\mathbf{x}'\mathbf{y}}{\mathbf{x}'\mathbf{x}} = \frac{\frac{1}{n}\mathbf{x}'\mathbf{y}}{\frac{1}{n}\|\mathbf{x}\|^2} = \frac{Kov(x, y)}{s_x^2} = \hat{b},$$

d.h. die aus der Statistik bekannte Schätzung für b . Auch $\tilde{b}$ ist eine Schätzung für b , auch wenn der Ausdruck für b direkt aus (8.18) folgt, denn die linke Seite von (8.19) enthält den Term $\mathbf{x}'\mathbf{e}/\|\mathbf{x}\|^2$, und der Fehlervektor $\mathbf{e}$ ist unbekannt; $\mathbf{x}'\mathbf{e}$ repräsentiert die Kovarianz von $\mathbf{x}$ und $\mathbf{e}$. Man kann nun einfach die *Annahme* machen, dass diese Kovarianz gleich Null ist und deswegen $\mathbf{x}'\mathbf{e} = 0$ setzen; dann entspricht die linke Seite gerade dem üblichen Ausdruck für den Regressionskoeffizienten. Man muß aber bedenken, dass $\mathbf{e}$ von Stichprobe zu Stichprobe variiert, so dass $\mathbf{x}'\mathbf{e}$ keinesfalls für eine *gegebene* Stichprobe gleich Null sein muß, sondern allenfalls im Mittel über die verschiedenen möglichen Stichproben gleich Null sein kann; das Postulat $\mathbb{E}(\mathbf{x}'\mathbf{e}) = 0$ bedeutet ja nicht, dass auch $\mathbf{x}'\mathbf{e} = 0$ ist.

Natürlich ist die aus der Statistik bekannte Schätzung $\hat{b}$ eine Kleinste-Quadrate-Schätzung. Nach der Methode der Kleinsten Quadrate wird b so geschätzt, dass die Summe der Quadrate der Fehler, also $\mathbf{e}'\mathbf{e}$, minimiert wird. Es soll also gelten

$$\mathbf{e}'\mathbf{e} = (\mathbf{y} - b\mathbf{x})'(\mathbf{y} - b\mathbf{x}) \stackrel{!}{=} \min. \quad (8.20)$$

Es ist

$$(\mathbf{y} - b\mathbf{x})'(\mathbf{y} - b\mathbf{x}) = \mathbf{y}'\mathbf{y} - b\mathbf{y}'\mathbf{x} - b\mathbf{x}'\mathbf{y} + b^2\mathbf{x}'\mathbf{x},$$

und differenziert man nach b und setzt die entstehende Ableitung gleich Null, so erhält man

$$-2\mathbf{x}'\mathbf{y} - 2\hat{b}\mathbf{x}'\mathbf{x} = 0,$$

woraus

$$\hat{b} = \frac{\mathbf{x}'\mathbf{y}}{\mathbf{x}'\mathbf{x}} = \frac{Kov(x, y)}{s_x^2} \quad (8.21)$$

folgt. Hier ist nur implizit angenommen worden, dass $\mathbf{e}$ von b unabhängig ist. Es kann nun leicht gezeigt werden, dass der Fehlervektor $\hat{\mathbf{e}} = \mathbf{y} - \hat{b}\mathbf{x}$ orthogonal zu $\mathbf{x}$ und $\hat{\mathbf{y}} = \hat{b}\mathbf{x}$ ist, so dass $\hat{\mathbf{e}}'\hat{\mathbf{y}} = \hat{\mathbf{e}}'\hat{b}\mathbf{x} = 0$. Denn es ist

$$\begin{aligned} \hat{\mathbf{e}}'\hat{\mathbf{y}} &= (\mathbf{y} - \hat{b}\mathbf{x})'\hat{b}\mathbf{x} \\ &= \hat{b}\mathbf{y}'\mathbf{x} - \hat{b}^2\mathbf{x}'\mathbf{x} = \hat{b}(\mathbf{x}'\mathbf{y} - \hat{b}\mathbf{x}'\mathbf{x}) \\ &= \hat{b}(\mathbf{x}'\mathbf{y} - \frac{\mathbf{x}'\mathbf{y}}{\mathbf{x}'\mathbf{x}}\mathbf{x}'\mathbf{x}) = 0, \end{aligned} \quad (8.22)$$

und dies gilt für alle Stichproben.

$\hat{\mathbf{y}}$ ist eine Linearkombination von $\mathbf{x}$, liegt also in dem durch $\mathbf{x}$ definierten 1-dimensionalen Teilraum des n -dimensionalen Vektorraums. $\mathbf{y}$ ist eine Linearkombination von $\hat{\mathbf{y}}$ und $\hat{\mathbf{e}}$:

$$\mathbf{y} = \hat{\mathbf{y}} + \mathbf{e}, \quad \hat{\mathbf{y}} \perp \mathbf{e}. \quad (8.23)$$

Die Orthogonalität von $\hat{\mathbf{y}}$ und $\mathbf{e}$ impliziert – wie man leicht nachrechnet –

$$\|\mathbf{y}\|^2 = \|\hat{\mathbf{y}}\|^2 + \|\mathbf{e}\|^2$$

und damit

$$1 = \frac{\|\hat{\mathbf{y}}\|^2}{\|\mathbf{y}\|^2} + \frac{\|\mathbf{e}\|^2}{\|\mathbf{y}\|^2},$$

und wegen

$$\frac{\|\hat{\mathbf{y}}\|^2}{\|\mathbf{y}\|^2} = \frac{\hat{b}^2 \|\mathbf{x}\|^2}{\|\mathbf{y}\|^2} = \frac{(\mathbf{x}'\mathbf{y})^2 \|\mathbf{x}\|^2}{\|\mathbf{x}\|^4 \|\mathbf{y}\|^2} = \frac{(\mathbf{x}'\mathbf{y})^2}{\|\mathbf{x}\|^2 \|\mathbf{y}\|^2} = r_{xy}^2$$

folgt

$$r_{xy}^2 = \frac{\|\hat{\mathbf{y}}\|^2}{\|\mathbf{y}\|^2} = \frac{s_{\hat{y}}^2}{s_y^2} = 1 - \frac{\|\mathbf{e}\|^2}{\|\mathbf{y}\|^2} = 1 - \frac{s_e^2}{s_y^2}, \quad (8.24)$$

bzw.

$$s_e^2 = s_y^2(1 - r_{xy}^2). \quad (8.25)$$

Projektion: $\hat{\mathbf{y}}$ kann als Projektion von $\mathbf{y}$ auf den durch $\mathbf{x}$ definierten Teilraum gesehen werden. Nach Abbildung 12, Seite 157, hat man (nach Vertauschung von $\mathbf{x}$ und $\mathbf{y}$)

$$\|\mathbf{p}_{yx}\| = \|\mathbf{y}\| \cos \theta = \|\mathbf{y}\| \frac{\mathbf{x}'\mathbf{y}}{\|\mathbf{x}\| \|\mathbf{y}\|} = \frac{\mathbf{x}'\mathbf{y}}{\|\mathbf{x}\|} = \hat{b} \|\mathbf{x}\| = \|\hat{\mathbf{y}}\|. \quad (8.26)$$

Diese Aussage gilt auch für die Verallgemeinerung der linearen Regression zur multiplen Regression.

8.5 Alternativer Beweis von Satz 3.27

Es sei $\mathbf{a}_{zi}$ der i -te Zeilenvektor von A , $i = 1, \dots, m$. Es sei $A\mathbf{x} = \vec{0}$; dies bedeutet, dass die Skalarprodukte $\mathbf{a}_{ui}'\mathbf{x} = 0$ für alle i , d.h. $\mathbf{x}$ ist orthogonal zu allen Zeilenvektoren von A . Für die $\mathbf{x}$ mit $A\mathbf{x} = \mathbf{y} \neq \vec{0}$ gilt diese Aussage nicht, d.h. $\mathbb{R}^n$ wird in zwei Teilmengen U und V zerlegt:

$$U = \{\mathbf{x} \in \mathbb{R}^n | A\mathbf{x} = \vec{0}\} = \text{kern}(A), \quad V = \{\mathbf{x} \in \mathbb{R}^n | A\mathbf{x} = \mathbf{y} \neq \vec{0}\}.$$

$U = \text{kern}(A)$ ist ein Teilraum des $\mathbb{R}^n$. V ist ebenfalls ein Teilraum des $\mathbb{R}^n$, denn es sei $A\mathbf{x}_1 = \mathbf{y}_1$, $A\mathbf{x}_2 = \mathbf{y}_2$, $\mathbf{x}_1, \mathbf{x}_2 \notin \text{kern}(A)$. Dann ist, für $\lambda, \mu \in \mathbb{R}$ beliebig, $A\lambda\mathbf{x}_1 = \lambda\mathbf{y}_1$, $A\mu\mathbf{x}_2 = \mu\mathbf{y}_2$ und $A(\lambda\mathbf{x}_1 + \mu\mathbf{x}_2) = \lambda\mathbf{y}_1 + \mu\mathbf{y}_2 = \mathbf{y} \in \mathcal{L}(A)$, mithin ist $\mathbf{x} = \lambda\mathbf{x}_1 + \mu\mathbf{x}_2 \in V$. Offenbar ist $U \cap V = \emptyset$, denn ein Vektor $\mathbf{x} \in \mathbb{R}^n$ kann nicht zugleich in U und in V sein. Also folgt $U + V = \mathbb{R}^n$ und nach dem Dimensionssatz 2.16, Seite 57, folgt

$$\dim(U + V) = \dim(U) + \dim(V) = \dim \mathbb{R}^n = n.$$

Zur Bestimmung von $\dim(U)$ werde

$$A\mathbf{x} = x_1\mathbf{a}_1 + \cdots + x_r\mathbf{a}_r + x_{r+1}\mathbf{a}_{r+1} + \cdots + x_n\mathbf{a}_n = \vec{0}$$

betrachtet. Ohne Beschränkung der Allgemeinheit kann angenommen werden, dass die Spaltenvektoren $\mathbf{a}_1, \dots, \mathbf{a}_r$ die linear unabhängigen Vektoren von A sind. Schreibt man

$$x_1\mathbf{a}_1 + \cdots + x_r\mathbf{a}_r = -(x_{r+1}\mathbf{a}_{r+1} + \cdots + x_n\mathbf{a}_n),$$

und berücksichtigt man, dass die $\mathbf{a}_{r+k}$ für $k = 1, \dots, n-r$ Linearkombinationen der $\mathbf{a}_j$, $j = 1, \dots, r$ sind, so dass

$$\mathbf{a}_{r+k} = \lambda_{k1}\mathbf{a}_1 + \cdots + \lambda_{kr}\mathbf{a}_r = \sum_{j=1}^r \lambda_{kj}\mathbf{a}_j \quad (8.27)$$

gelten muß, so hat man

$$\sum_{j=1}^r x_j\mathbf{a}_j = - \sum_{k=1}^{n-r} x_{r+k} \sum_{j=1}^r \lambda_{kj}\mathbf{a}_j. \quad (8.28)$$

Über den Vektor $\mathbf{x}$ ist bisher keine weitere Annahme gemacht worden außer, dass $A\mathbf{x} = \vec{0}$ gelten soll. Man kann also insbesondere

$$\mathbf{x} = \mathbf{x}_{r+k} = (x_1, \dots, x_r, 0, \dots, 0, 1, 0, \dots, 0)', \quad k = 1, \dots, n-r$$

setzen, wobei die 1 an der $(r+k)$ -ten Stelle stehen soll, also

$$\begin{aligned} \mathbf{x}_{r+1} &= (x_1, \dots, x_r, 1, 0, \dots, 0)', \\ \mathbf{x}_{r+2} &= (x_1, \dots, x_r, 0, 1, 0, \dots, 0)', \\ \mathbf{x}_{r+3} &= (x_1, \dots, x_r, 0, 0, 1, 0, \dots, 0)' \\ &\vdots \end{aligned}$$

Die Gleichung (8.28) nimmt dann die Form

$$\sum_{k=1}^r x_{r+k} \sum_{j=1}^r \mu_{kj}\mathbf{a}_j = \sum_{j=1}^r \lambda_{kj}\mathbf{a}_j = - \sum_{j=1}^r x_j\mathbf{a}_j.$$

so dass

$$\sum_{j=1}^r \lambda_{kj}\mathbf{a}_j + \sum_{j=1}^r x_j\mathbf{a}_j = \sum_{j=1}^n (x_j + \lambda_j)\mathbf{a}_j = 0.$$

Da die $\mathbf{a}_j$ linear unabhängig sind folgt $x_j + \lambda_j = 0$ oder $-\lambda_j = x_j$ für alle j . es ist also

$$\mathbf{x}_k = (-\lambda_{k1}, \dots, -\lambda_{kr}, 0, \dots, 0, 1, 0, \dots, 0)', \quad k = 1, \dots, n-r \quad (8.29)$$

wobei die 1 an der $(r+k)$ -ten Stelle steht. In der Tat ist

$$A\mathbf{x}_k = -\lambda_{k1}\mathbf{a}_1 - \cdots - \lambda_{kr} + \mathbf{a}_{r+k} = 0$$

wegen (8.27).

Die $\mathbf{x}_k$ sind linear unabhängig, wie man sofort sieht, denn

$$\mu_1\mathbf{x}_1 + \cdots + \mu_k\mathbf{x}_k = 0$$

impliziert $\mu_i = 0$ für $i = 1, \dots, n-r$. Sie bilden damit eine Basis für einen $(n-r)$ -dimensionalen Teilraum. Damit ist gezeigt, dass $\text{kern}(A)$ mindestens $(n-r)$ -dimensional ist. Die Frage ist, ob die Dimensionalität von $\text{kern}(A)$ nicht größer ist. Das ist aber nicht möglich, da ja bereits $\text{rg}(A) = r$ angenommen wurde, der Rang von $\text{kern}(A)$ kann also nicht größer als $n-r$ sein. Wegen (??) (= Dimensionssatz) folgt weiter, dass $\text{rg}(V) = \mathcal{L}(A) = r$ ist. $\square$

Für $r = n$ folgt demnach $\text{rg}[\text{kern}(A)] = 0$, d.h. in diesem Fall hat $A\mathbf{x} = \vec{0}$ nur eine Lösung: $\mathbf{x} = \vec{0}$. Dies ist evident, denn in diesem Fall sind die $\mathbf{a}_1, \dots, \mathbf{a}_n$ linear unabhängig und $\sum_j x_j \mathbf{a}_j = \vec{0}$ nur dann, wenn $x_1 = \cdots = x_n = 0$. $\square$

8.6 Gleichungssysteme und Singularwertzerlegung (II)

Die Koeffizientenmatrix sei A eine $m \times n$ -Matrix, $m \geq n$. Die Singularwertzerlegung sei durch

$$A = Q\Sigma P', \quad \Sigma = \Lambda^{1/2} \quad (8.30)$$

mit $\sigma_j = \sqrt{\lambda_j}$ gegeben, wobei λ_j der j -te Eigenwert von $A'A$ bzw. AA' ist. $r \leq n$ sei der Rang von A . Ist $r = n$, so sind n Eigenwerte ungleich Null, ist dagegen $r < n$, so sind nur r Eigenwerte ungleich Null und Σ hat die Form

$$\Sigma = \begin{pmatrix} \sigma_1 & 0 & 0 & 0 & 0 & \cdots & 0 \\ 0 & \sigma_2 & 0 & 0 & 0 & \cdots & 0 \\ 0 & 0 & \ddots & 0 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & \sigma_r & 0 & \cdots & 0 \\ 0 & 0 & \cdots & 0 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & 0 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & 0 & 0 & \cdots & 0 \end{pmatrix} = \begin{pmatrix} \Sigma_r & 0 \\ 0 & 0 \end{pmatrix}, \quad \sigma_j = \sqrt{\lambda_j} \quad (8.31)$$

d.h. Σ ist eine $n \times n$ -Matrix, deren erste r Diagonalelemente von Null verschieden sind und bei der letzten $n-r$ Zeilen und Spalten nur Nullen enthalten; σ_j heißt j -ter *Singularwert*. $\Sigma_r = \Lambda_r^{1/2}$ ist eine Diagonalmatrix, die nur die von Null verschiedenen Singularwerte enthält.

Für den Fall $r < n$ sind $\sigma_1, \dots, \sigma_r \neq 0$, und $\sigma_{r+1} = \cdots = \sigma_n = 0$. Σ ist eine $n \times n$ -Matrix, deren letzte $n-r$ Spalten ebenso wie die letzten $n-r$ Zeilen nur Nullen enthält. P ist eine $n \times n$ -Matrix orthonormaler Vektoren (die Eigenvektoren

von $A'A$). Die letzten $n-r$ (Spalten-)Vektoren $\mathbf{P}_{r+1}, \dots, \mathbf{P}_n$ von P korrespondieren zu Eigenwerten und damit zu σ -Werten, die gleich Null sind. Es sei r der Rang von A und $\mathcal{B}_r = \{\mathbf{P}_1, \dots, \mathbf{P}_r\}$; $\mathcal{L}(\mathcal{B}_r)$ sei die Menge der n -dimensionalen Vektoren, die sich als Linearkombination der $\mathbf{P}_1, \dots, \mathbf{P}_r$ darstellen lassen. Weiter sei $\mathcal{B}_{n-r} = \{\mathbf{P}_{r+1}, \dots, \mathbf{P}_n\}$ und $\mathcal{L}(\mathcal{B}_{n-r})$ die Menge der n -dimensionalen Vektoren, die sich als Linearkombinationen der $\mathbf{P}_{r+1}, \dots, \mathbf{P}_n$ darstellen lassen. $\mathcal{L}(\mathcal{B}_r)$ ist ein r -dimensionaler Teilraum des V_n , und $\mathcal{L}(\mathcal{B}_{n-r})$ ist ein $n-r$ -dimensionaler Teilraum des V_n . $\mathcal{B}_r \cup \mathcal{B}_{n-r}$ bildet eine Basis für den gesamten V_n . Für den Fall $n=r$ ist $\mathcal{B}_{n-r}$ leer.

Satz 8.3 *Es sei A eine $m \times n$ -Matrix mit $m \geq n$ und Rang $r \leq n$; die SVD von A sei durch (8.30) gegeben. Es gelte*

$$A\mathbf{x} = \mathbf{y}, \quad (8.32)$$

wobei $\mathbf{x}$ und $\mathbf{y}$ n -dimensionale Vektoren sind. Dann gilt

$$\mathbf{y} = \vec{0} \Rightarrow \mathbf{x} = a_1 \mathbf{P}_{r+1} + \dots + a_{n-r} \mathbf{P}_n \in \mathcal{L}(\mathcal{B}_{n-r}) \quad (8.33)$$

$$\mathbf{y} \neq \vec{0} \Rightarrow \mathbf{x} = b_1 \mathbf{P}_1 + \dots + b_r \mathbf{P}_r \in \mathcal{L}(\mathcal{B}_r), \quad (8.34)$$

wobei die $\mathbf{P}_1, \dots, \mathbf{P}_r, \mathbf{P}_{r+1}, \dots, \mathbf{P}_n$ die Spaltenvektoren von P in (8.30) sind und die $a_1, \dots, a_{n-r}$ bzw. die $b_1, \dots, b_r$ geeignet gewählte Koeffizienten sind.

Beweis: Es sei $\mathbf{y} = \vec{0}$ und es gelte $r < n$. Da P orthonormal ist, bilden die Spaltenvektoren $\mathbf{P}_1, \dots, \mathbf{P}_n$ eine orthonormale Basis des V_n , und da $\mathbf{x} \in V_n$ kann $\mathbf{x}$ stets als Linearkombination

$$\mathbf{x} = c_1 \mathbf{P}_1 + \dots + c_n \mathbf{P}_n \quad (8.35)$$

dargestellt werden. Es ist zu zeigen, dass für $A\mathbf{x} = \vec{0}$ stets $c_1 = \dots = c_r = 0$ gilt, und für $A\mathbf{x} \neq \mathbf{y}$ folgt $c_{r+1} = \dots = c_n = 0$.

Es gelte (8.33), d.h. es gelte $A\mathbf{x} = \vec{0}$. (8.35) kann in der Form

$$\mathbf{x} = \sum_{j=1}^r c_j \mathbf{P}_j + \sum_{k=r+1}^n c_k \mathbf{P}_k \quad (8.36)$$

geschrieben werden. Es sei $A\mathbf{x} = Q\Sigma P'\mathbf{x} = \vec{0}$, d.h.

$$\vec{0} = Q\Sigma \sum_{j=1}^r c_j \mathbf{P}_j + Q\Sigma \sum_{k=r+1}^n c_k \mathbf{P}_k.$$

Es ist aber

$$\begin{aligned} Q\Sigma P' \sum_{k=r+1}^n c_k \mathbf{P}_k &= Q\Sigma \sum_{k=r+1}^n c_k P' \mathbf{P}_k \\ &= Q \sum_{k=r+1}^n c_k \Sigma \mathbf{e}_k = Q \sum_{k=r+1}^n c_k \mathbf{s}_k \end{aligned} \quad (8.37)$$

wobei $\mathbf{e}_k = P' \mathbf{P}_k$ der k -te Einheitsvektor ist (seine Komponenten sind alle gleich Null bis auf die k -te, die gleich 1 ist), denn die Vektoren in P sind orthonormal, und $\mathbf{s}_k = \Sigma \mathbf{e}_k$ ist der k -te Spaltenvektor von Σ ⁴⁷. Aber für $k > r$ ist dieser Spaltenvektor der Nullvektor (für $r < n$ gibt es nur r Eigenwerte ungleich Null, vergl. (8.31)), so dass unabhängig von der Wahl der Koeffizienten c_k die (Teil-)Summe $Q \Sigma P' \sum_{k=r+1}^n c_k \mathbf{P}_k$ gleich Null ist.

Nun werde die erste Teilsumme $\sum_{j=1}^r c_j \mathbf{P}_j$ für $\mathbf{x}$ betrachtet. Da die zweite Teilsumme auf jeden Fall gleich Null ist, muß nun

$$A\mathbf{x} = Q \Sigma P' \mathbf{x} = Q \Sigma P' \sum_{j=1}^r c_j \mathbf{P}_j = \vec{0}$$

gelten. Es ist aber

$$\begin{aligned} Q \Sigma P' \sum_{j=1}^r c_j \mathbf{P}_j &= Q \Sigma \sum_{j=1}^r c_j P' \mathbf{P}_j \\ &= Q \Sigma \sum_{j=1}^r c_j \mathbf{e}_j = Q \sum_{j=1}^r c_j \mathbf{s}_j \end{aligned} \quad (8.38)$$

wobei $\mathbf{s}_j$ der j -te Spaltenvektor von Σ ist. Für $j \leq r$ ist aber $\mathbf{s}_j \neq \vec{0}$, so dass $Q \Sigma P' \sum_{j=1}^r c_j \mathbf{P}_j \neq \vec{0}$ ist, es sei denn $c_1 = \dots = c_r = 0$. Damit ist gezeigt, dass für den Fall $A\mathbf{x} = \vec{0}$ und $n < r$ der Vektor keine Linearkombination der $\mathbf{P}_1, \dots, \mathbf{P}_r$ sein kann, aber als Linearkombination der $\mathbf{P}_{r+1}, \dots, \mathbf{P}_n$ dargestellt werden kann. $\square$

Anmerkungen:

1. Für den Fall $r = n$ folgt sofort, dass $\mathcal{B}_{n-r} = \mathcal{B}_0 = \emptyset$, d.h. es existiert kein Vektor $\mathbf{x} \neq \vec{0}$ derart, dass $A\mathbf{x} = \vec{0}$ ist.
2. Da $\mathbf{x} = a_1 \mathbf{P}_{r+1} + \dots + a_{n-r} \mathbf{P}_n$ die Gleichung $A\mathbf{x} = \vec{0}$ impliziert, folgt aus $A\mathbf{x} \neq \vec{0}$, dass $\mathbf{x} \notin \mathcal{B}_{n-r}$, d.h. $\mathbf{x}$ kann kein Element des aus den $\mathbf{P}_{r+1}, \dots, \mathbf{P}_n$ aufgespannten Teilraums sein. Da aber $\mathbf{v} \in V_n$, muß $\mathbf{x}$ ein Element des aus den $\mathbf{P}_1, \dots, \mathbf{P}_r$ aufgespannten Teilraums sein, d.h. $\mathbf{x}$ ist als Linearkombination der $\mathbf{P}_1, \dots, \mathbf{P}_r$ darstellbar.
3. Mit $A\mathbf{x} = \mathbf{y}$ und $A = Q \Sigma P'$ gilt $Q \Sigma P' \mathbf{x} = \mathbf{y}$ und man erhält

$$\mathbf{x} = P \Sigma^{-1} Q' \mathbf{y}, \quad (8.39)$$

d.h. die Lösung für $\mathbf{x}$ läßt sich über die SVD von A bestimmen. Ist $r < n$, so kann man $A = Q_r \Sigma_r P_r'$ schreiben, wobei Σ_r eine $r \times r$ -Diagonalmatrix

⁴⁷Es sei S eine beliebige $n \times n$ -Matrix; dann ist $S \mathbf{e}_k = \mathbf{s}_k$ stets der k -te Spaltenvektor von S , wie man leicht nachrechnet

ist, die nur die von Null verschiedenen Singularwerte (= Wurzeln aus den Eigenwerten) enthält. Man hat dann eine Lösung

$$\mathbf{x} = P_r \Sigma_r^{-1} Q_r' \mathbf{y}. \quad (8.40)$$

Die Lösung (8.39) entspricht natürlich der Lösung (3.296), wie man sieht, wenn man in (3.296) A durch $Q \Sigma P'$ oder im Fall $r < n$ A durch $Q_r \Sigma_r P_r'$ ersetzt. Man vergleiche die Lösungen mit denen für die verallgemeinerte Inverse in Abschnitt 3.15.

8.7 Alternative Herleitung der PCA

Gesucht sind linear unabhängige, insbesondere orthogonale Vektoren $\mathbf{L}_1, \dots, \mathbf{L}_n$, aus denen sich die Spaltenvektoren einer Datenmatrix X als Linearkombinationen darstellen lassen. Da nur X gegeben ist, müssen sich die $\mathbf{L}_k$ aus X errechnen lassen. Insbesondere muß sich $\mathbf{L}_1$ aus X durch Multiplikation von X mit einem Vektor $\mathbf{t}_1$ bestimmen lassen: $X \mathbf{t}_1 = \mathbf{L}_1$. Unter geeigneten Normierungsbedingungen ist $\mathbf{L}_1' \mathbf{L}_1$ proportional (dh bis auf einen Faktor $1/m$) zur Varianz der Komponenten von $\mathbf{L}_1$. $\mathbf{t}_1$ soll so bestimmt werden, dass $\mathbf{L}_1' \mathbf{L}_1$ maximal wird. Ohne weitere Nebenbedingung wird diese Varianz aber maximal, wenn $\mathbf{L}_1' \mathbf{L}_1 = \infty$ wird, was nicht sinnvoll ist, die Komponenten von $\mathbf{L}_1$ haben dann keine endlichen Werte mehr. Deswegen wird die Nebenbedingung $\mathbf{t}_1' \mathbf{t}_1 = 1$ eingeführt; diese Bedingung definiert $\phi(t_{11}, \dots, t_{n1}) = 0$. Also soll gelten

$$Q(\mathbf{L}) = \mathbf{t}_1' X' X \mathbf{t}_1 + \lambda(\mathbf{t}_1' \mathbf{t}_1 - 1) = \mathbf{L}_1' \mathbf{L}_1 = \max \quad (8.41)$$

Man muß diese Gleichung nach den Komponenten t_{i1} und nach λ differenzieren und die entstehenden Gleichungen gleich Null setzen. Man hat bezüglich t_{i1}

$$\frac{\partial L}{\partial t_{i1}} = \mathbf{e}_i' X' X \mathbf{t}_1 + \mathbf{t}_1' X' X \mathbf{e}_i + \lambda(\mathbf{e}_i' \mathbf{t}_1 + \mathbf{t}_1' \mathbf{e}_i) = 0,$$

Schreibt man dies für alle $\mathbf{e}_i$ und berücksichtigt, dass $\mathbf{e}_i' X' X \mathbf{t}_1$, $\mathbf{t}_1' X' X \mathbf{e}_i$, $\mathbf{e}_i' \mathbf{t}_1$ und $\mathbf{t}_1' \mathbf{e}_i$ Skalarprodukte mit identischem Wert sind, so erhält man mit $\lambda_1 = \lambda$ die Gleichung

$$X' X \mathbf{t}_1 = \lambda_1 \mathbf{t}_1, \quad (8.42)$$

d.h. $\mathbf{t}_1$ ist der erste Eigenvektor von $X' X$ und λ_1 ist der zugehörige Eigenwert.

Zur Bestimmung von $\mathbf{L}_2$ verfährt man analog; man hat aber wegen der Orthogonalitätsforderung die Nebenbedingungen $\mathbf{t}_2' \mathbf{t}_2 = 1$ und $\mathbf{L}_1' \mathbf{L}_2 = 0$. Dies bedeutet

$$\mathbf{L}_2' \mathbf{L}_2 = \mathbf{t}_2' X' X \mathbf{t}_2 + \lambda_2(\mathbf{t}_2' \mathbf{t}_2 - 1) + \mu(\mathbf{t}_2' X' X \mathbf{t}_1 - 0)$$

soll maximiert werden. Differenziert man wie im ersten Fall und setzt die entstehenden Ableitungen gleich Null, erhält man

$$X' X \mathbf{t}_2 = \lambda_2 \mathbf{t}_2, \quad \mu = 0, \quad (8.43)$$

d.h. $\mathbf{L}_2$ wird durch den zweiten Eigenvektor $\mathbf{t}_2$ von $X'X$ bestimmt, mit dem zugehörigen Eigenwert λ_2 . So verfährt man weiter und erhält die Transformation

$$XT = L, \quad T'T = I, \quad (8.44)$$

und diese Lösung entspricht der Lösung $XP = L$, auf die man geführt wird, wenn man von $X = LP'$ zusammen mit der Orthogonalitätsforderung $L'L = \Lambda$ ausgeht.

8.8 Ein Maximum-Prinzip

In Gleichung (2.55) wurde die Cauchy-Schwarzsche Ungleichung

$$|\mathbf{x}'\mathbf{y}| \leq \|\mathbf{x}\| \|\mathbf{y}\|$$

eingeführt. Diese Ungleichung kann verallgemeinert werden.

Es sei M eine positiv-definite $(n \times n)$ -Matrix und $\mathbf{x}$ und $\mathbf{y}$ seien zwei n -dimensionale Vektoren. Dann gilt

$$(\mathbf{x}'\mathbf{y})^2 \leq (\mathbf{x}'M\mathbf{x})(\mathbf{y}'M^{-1}\mathbf{y}). \quad (8.45)$$

Beweis: Für $\mathbf{x} = \vec{0}$ und/oder $\mathbf{y} = \vec{0}$ ist die Ungleichung trivialerweise wahr, $\mathbf{x}$ und $\mathbf{y}$ seien also vom Nullvektor verschieden. Es sei

$$M^{1/2} = \sum_{k=1}^n \sqrt{\lambda_k} \mathbf{P}_k \mathbf{P}'_k, \quad M^{-1/2} = \sum_{k=1}^n \frac{\mathbf{P}_k \mathbf{P}'_k}{\sqrt{\lambda_k}}.$$

Dann folgt

$$\mathbf{x}'\mathbf{y} = \mathbf{x}'I\mathbf{y} = \mathbf{x}'M^{1/2}M^{-1/2}\mathbf{y} = (M^{1/2}\mathbf{x})'(M^{-1/2}\mathbf{y}).$$

$M^{1/2}\mathbf{x}$ und $M^{-1/2}\mathbf{y}$ sind aber Vektoren, für die die Cauchy-Schwarzsche Ungleichung existiert, dh

$$|(M^{1/2}\mathbf{x})'(M^{-1/2}\mathbf{y})| = |\mathbf{x}'M^{1/2}M^{-1}\mathbf{y}| = |\mathbf{x}'\mathbf{y}| \leq |M^{1/2}\mathbf{x}| |M^{-1/2}\mathbf{y}|,$$

und das war zu zeigen. $\square$

Das Maximum-Prinzip Es sei M wieder eine positiv-definite $(n \times n)$ -Matrix und $\mathbf{x}$ sei ein n -dimensionaler Vektor. Ebenso sei $\mathbf{y} \neq \vec{0}$ ein n -dimensionaler Vektor. Dann gilt

$$\max_{\mathbf{y} \neq \vec{0}} \frac{(\mathbf{y}'\mathbf{x})^2}{\mathbf{y}'M\mathbf{y}} = \mathbf{x}'M^{-1}\mathbf{x}. \quad (8.46)$$

Das Maximum wird angenommen für $\mathbf{y} = cM^{-1}\mathbf{x}$, für ein $c \in \mathbb{R}$.

Beweis: Nach (8.45) gilt

$$\mathbf{x}'\mathbf{y} \leq (\mathbf{y}'M\mathbf{y})(\mathbf{x}'M^{-1}\mathbf{x}).$$

Da M positiv-definit ist, gilt $\mathbf{y}'M\mathbf{y} > 0$. Dividiert man beide Seiten der Ungleichung durch $\mathbf{y}'M\mathbf{y}$, erhält man

$$\frac{(\mathbf{x}'\mathbf{y})^2}{\mathbf{y}'M\mathbf{y}} \leq \mathbf{x}'M^{-1}\mathbf{x}.$$

Für $\mathbf{y} = cM^{-1}\mathbf{x}$ wird das Maximum angenommen. □

8.9 Die n -dimensionale Normalverteilung

Sie ist durch

$$f(\mathbf{x}) = A \exp(-(\mathbf{x} - \mu)' \Sigma^{-1} (\mathbf{x} - \mu)) \quad (8.47)$$

definiert, wobei $\mathbf{x}$ ein n -dimensionaler Vektor ist, μ ist ebenfalls ein n -dimensionaler Vektor, dessen Komponenten die Mittelwerte der Komponenten von $\mathbf{x}$ sind, und Σ ist die Varianz-Kovarianzmatrix der Komponenten von $\mathbf{x}$. A ist ein Normierungsfaktor, der bewirkt, dass $\int f(\mathbf{x}) d\mathbf{x} = 1$. Offenbar ist die Dichte durch quadratische Formen definiert: $f(\mathbf{x})$ für $(\mathbf{x} - \mu)' \Sigma^{-1} (\mathbf{x} - \mu) = k$ ist die Dichte für alle $\mathbf{x}$, deren Endpunkte auf einer bestimmten Ellipse liegen, deren Hauptachsenlängen durch den Wert von k (vergl. (3.138) und (3.139) bestimmt werden. Σ^{-1} ist symmetrisch, da Σ symmetrisch ist. Wegen (3.157) gilt dann auch

$$(\mathbf{x} - \mu)' \Sigma^{-1} (\mathbf{x} - \mu) = (\mathbf{x} - \mu)' \left(\sum_{k=1}^n \frac{\mathbf{t}_k \mathbf{t}_k'}{\lambda_k} \right) (\mathbf{x} - \mu), \quad (8.48)$$

wobei die $\mathbf{t}_k$ die Eigenvektoren von Σ sind und λ_k die entsprechenden Eigenwerte. Die Elemente von Σ^{-1} sind offenbar groß, wenn es kleine Eigenwerte gibt.

Im Übrigen bedeutet die Zerlegung der Korrelationsmatrix R gemäß $R = T\Lambda T'$ nicht, dass implizit die Normalverteilung angenommen wird. Die Zerlegung ist einfach eine Implikation der Symmetrie von R .

Mahalanobis-Distanz Der Ausdruck

$$\delta^2 = (\mathbf{x} - \mu)' \Sigma^{-1} (\mathbf{x} - \mu) \quad (8.49)$$

heißt *Mahalanobis-Distanz*. Alle $\mathbf{x}$ mit $\delta^2 = \text{konstant}$ haben die gleiche Dichte, d.h. die gleiche Wahrscheinlichkeit, zwischen $\mathbf{x}$ und $\mathbf{x} + d\mathbf{x}$ zu liegen. Eine *euklidisch* größere Distanz in Richtung der ersten Hauptachse hat die gleiche Wahrscheinlichkeit wie eine euklidisch kleinere Distanz in Richtung der zweiten Hauptachse (und analog für die weiteren Hauptachsen des Ellipsoids). Der Begriff der Entfernung (Distanz) wird hier mit dem der Wahrscheinlichkeit kombiniert.

Literatur

- [1] Aronszajn, N. (1950) Theory of reproducing kernels. *Trans. Amer. Math. Soc.*, 68, 337–404
- [2] Barrantes Campos, H.: Elementos de álgebra lineal. Editoria Universidad Estatal a Distancia, San José, Costa Rica 2012
- [3] Basilevsky, A.: Statistical Factor Analysis and Related Methods – Theory and Applications. John Wiley & Sons, New York etc 1994
- [4] Bedürftig, T., Murawski, R.: Philosophie der Mathematik. Berlin Boston 2012
- [5] Cadima J., Jolliffe, I. (2009) On relationships between uncentered and column-centered principal component analysis. *Pakistan Journal of Statistics*, 25 (4), 472 – 503
- [6] Courant, R., Hilbert, D.: Methoden der mathematischen Physik I. Springer-Verlag Berlin, Heidelberg, New York 1968
- [7] Eckart, C., Young, G. (1936), The approximation of one matrix by another of lower rank. *Psychometrika*, 1 (3): 211–8. doi:10.1007/BF02288367.
- [8] Goldstein, H.: Klassische Mechanik. Wiesbaden 1981
- [9] Fischer, G.: Lineare Algebra. Friedr. Vieweg & Sohn Verlagsgesellschaft, Braunschweig Wiesbaden 1997
- [10] Jolliffe, I.T.: Principal Component Analysis. Springer-Verlag, Berlin 1986
- [11] Koecher, M.: Lineare Algebra und analytische Geometrie. Berline, Heidelberg 1997
- [12] Lorenz, F.: Lineare Algebra I, BI Wissenschaftsverlag, Mannheim/Zürich 1988
- [13] Lorenz, F.: Lineare Algebra II, BI Wissenschaftsverlag, Mannheim/Zürich 1988
- [14] Meschkowski, H.: Hilbertsche Räume mit Kernfunktionen. Springer-Verlag Berlin, Göttingen Heidelberg 1962
- [15] Mirsky, L. (1960) Symmetric gauge functions and unitary invariant norms. *Quarterly J. Math. Oxford*, 11, 50–590
- [16] Moore, E. H. (1920) On the reciprocal of the general algebraic matrix. *Bulletin of the American Mathematical Society*, 26, 394–395

- [17] Penrose, R. (1954) A generalized inverse for matrices. *Proceedings of the Cambridge Philosophical Society*, 51, 406–413
- [18] Schmidt, E. (1907) Zur Theorie der linearen und nichtlinearen Integralgleichungen. 1. Teil: Entwicklung willkürlicher Funktionen nach Systemen vorgeschriebener. *Mathematische Annalen*, 63, 433–476
- [19] Schölkopf, B., Smola, A.J.: Learning with kernels. The MIT Press, Cambridge (Massachusetts), London 2002

Index

- L^p -Räume, 191
- ℓ^p -Folgen, 191
- Ähnlichkeit, 25
- äußeres Produkt, 29
- Abbildung
 - bijektiv, 71, 167
 - Bild von, 70
 - Hintereinanderschaltung, 167
 - homomorphe, 71
 - identische, 70
 - injektiv, 71, 166
 - inverse, 70
 - isomorphe, 72
 - Kern einer, 71
 - Komposition, 167
 - lineare, 71
 - Rang der Abbildung, 72
 - surjektiv, 71, 166
 - Verknüpfung, 167
- Abbildung eines Vektors, 65
- abgeschlossene Teilmenge, 200
- achsenparallel, 98
- Adjunkte einer Matrix, 93
- affin, 18
- Algorithmus, Gauß-, 153
- Allgemeines lin. Modell, 122
- Anfangsabbrechfehler, 165
- Austauschsatz (Steiner), 51
- Automorphismus, 72
- Automorphismus, 170
- Banachraum, 193
- Basis
 - eines Vektorraums, 47
 - kanonische, 48
 - orthonormale, 48, 49
- Basisentwicklung eines Vektors
 - orthonormale, 48
- Basisfunktionen, 197
- Basisvektoren, 47
- Besselsche Ungleichung, 198
- Bild einer Abbildung, 172
- Cauchy-Folge, 190
- Cauchy-Schwarzsche Ungleichung, 28
- charakteristisches Polynom, 109
- Cholesky-Zerlegung, 155
- Cramersche Regel, 84, 152
- Defekt einer Abbildung, 172
- Deflation einer Matrix, 146
- Determinante, 83
- diagonalisierbar, 182
- Diagonalisierung, 110
- Dimension, 50
- Dimension Vektorraum, 51
- Dimensionssatz, 57
- Distanz
 - euklidische, 237
 - Mahalanobis-, 237
- Dreiecksmatrix
 - obere, 62
 - untere, 62
- dyadisches Produkt, 29
- Ebenen im V_n , 39
- Ebenengleichung, 39, 214
- Eigenfunktion, 205
- Eigenraum, 110, 187
- Eigenvektor, 99
 - Links-, 180
 - Rechts-, 180
- Eigenvektoren, 95
- Eigenwert, 99, 205
 - komplexer, 183
 - Nullstellen eines Polynoms, 183
 - Varianz der $\mathbf{L}_k$ -Komponenten, 141
- Einheitsmatrix, 62
- Einheitsvektor, 15
- Einsvektor, 15, 24
- elementare Umformung, 56, 81
- Eliminationsverfahren, gaußsches, 153
- Endomorphismus, 72, 110, 170

Epimorphismus, 170
 Erzeugendensystem, 50
 Faser, 71
 Fredholm-Typ, 207
 Fundamental-Lemma, 53
 Funktional, 202
 Gleichung, charakteristische, 179
 Gradientenvektor, 126
 Grundmengen, 200
 Häufungsstellenprinzip, Weierstraßsches, 193
 hat-matrix, 161
 Hauptachsentransformation, 107, 141
 Hauptraum, 188
 Hauptvektor, 188
 Hilbertraum
 separierbarer, 201
 Homomorphismus, 72, 169
 Hyperebene, 17, 41
 idempotent, 69
 Identität, 70
 Inverse Matrix, 89
 Isomorphismus, 72, 170
 Iteration, 147
 Körper, 225
 Kern, 149
 einer Abbildung, 172
 reproduzierender, 209
 Kernfunktion, 203, 216
 Kleinste Quadrate, 35
 Entwicklung von Funktionen, 198
 generalisierte, 124
 Methode der, 122
 Kodimension, 50
 Kofaktor, 85, 88, 93
 kollinear, 11, 34, 35
 Kollinearität, 163
 kompakt
 relativ, 204
 Konditionierungszahl, 164
 konkave Funktion, 223
 Konvergenz, 190
 gleichmäßige, 200
 konvex, 223
 konvexe Funktion, 192
 Koordinaten eines Vektors bezüglich einer Basis, 47, 49
 Koordinatenabbildung, 171
 Korrelation
 kanonische, 184
 Korrelationskoeffizient, 24
 Kosinussatz, 21
 Kroneckerprodukt, 165
 längeninvariant, 95
 Lösungsvektor, 33
 Ladung einer Variablen, 144
 Lagrange-Faktor, 126
 Lagrange-Funktion, 126
 Laplacescher Entwicklungssatz, 89
 LDL-Zerlegung, 156
 linear unabhängig, 197
 lineare Hülle, 45
 lineare Hülle (Gleich'syst.), 149
 Linearkombination, 18, 197
 Linksinverse, 95
 Linkssystem, 222
 Matrix
 adjungierte, 88
 assoziierte, 183
 charakteristische Gleichung, 179
 Diagonal-, 62
 gestürzte, transponierte, 61
 hermitesch, 63, 184
 idempotent, 69, 123
 imaginär, 63, 183
 inverse, 89
 konjugiert komplexe, 63
 konjugierte, 183
 Präzisions-, 93
 Projektions-, 161
 reell, 63
 schief-hermitesch, 63

- schief-symmetrisch, 63, 184
 - Skalierungs-, 62, 67
 - Spektraldarstellung, 110
 - Sub-, 63
 - symmetrische, 62
 - Wurzel aus einer Matrix, 114
- Matrixnorm, 131
- Matrizen
 - ähnliche, 182
- maximal linear unabhängig, 46
- Metrik, 12
- Metrik, euklidische, 12
- Minkowski-Metrik, 12
- Minor, 93
- Mittelwert, 24
- Modell, 27
- Monomorphismus, 170
- Multikollinearität, 163
- negativ semidefinit, 97
- Norm, 21
 - p -, 132
 - Matrix, 131
 - eines Operators, 203
 - euklidische, 131
 - Frobenius-, 132
 - Hilbert-Schmidt-, 132
 - induzierte, 132
 - Maximum, 131
 - Schur-, 132
- Normalenvektor, 20, 39, 172
- Nullmatrix, 62
- Nullvektor, 15
- Operator, 202
 - adjungierter, 203
 - Integral-, 203
 - kompakter, 205
- Orientierungsinvarianz, 95
- orthogonal, 23
- orthogonales Komplement, 55
- Orthonormalbasis (ONB), 48
- orthonormale Basisentwicklung, 48
- orthonormales Basissystem, 197
- Ortsvektor, 15
- Pivot-Element, 81
- Polynom, charakteristisches, 179
- positiv semidefinit, 97, 102
- Prähilbertraum, 195
- Präzisionsmatrix, 93
- Produkt
 - Kreuz, 20
 - äußeres, 20
 - äußeres (dyadisches), 29
 - dyadisches, 29, 111, 134
 - inneres, 19
 - Skalar-, 19
 - vektorielles, 20
- Projektion
 - Vektor auf einen anderen, 157
 - Vektor auf Teilraum, 230
- Projektionsmatrix, 161
- Projektionsoperator, 161
- Punktspektrum, 207
- quadratische Form, 97
- Rang
 - einer Matrix, 81
 - einer Vektormenge, 49, 56
 - voller, 49
- Raum, metrischer, 190
- Rayleigh-Quotient, 127
 - generalisierter, 184
- Rechtsinverse, 95
- Rechtssystem, 222
- Regression
 - multiple, 26
- Repräsentant eines Vektors, 14
- Resolventenmenge, 207
- Richtungsvektor, 16
- Rotation, 97
- Rotationsmatrix, 95
- Satz
 - von Courant-Fischer, 127
 - Eckart-Young, 135
 - Mercer, 211

- Pythagoras, 198
- Riesz, 210
- von Schmidt-Mirsky, 135, 136
- Schur-Norm, 132
- Scree-Test, 146
- Sensitivität, 164
- separierbar, 201
- Singularvektoren, 117
- Singularwert, 232
- Singularwerte, 117
- Singularwertzerlegung, 116, 117
- Skalar, 9
- Skalarprodukt
 - als Ähnlichkeitsmaß, 25
 - Eigenschaften-, 21
- Spaltenrang, 78
- Spaltenraum, 73
- Spann (span) einer Matrix, 45
- Spektraldarstellung, 101, 111
- Spektralsatz, 206
- Spur, 64
- stetig
 - gleichgradig, 195
 - Operator, 202
- Submatrizen, 63
- support vectors, 214, 216
- Teilbasis, 49
- Teilraum
 - invarianter, 188
- Teilvektorraum, 38
- Testmodell
 - faktorenanalytisches, 27
- Transformation eines Vektors, 65
- Umformungen, elementare, 153
- unendlich dimensional, 197
- Ungleichung
 - arithm. und geom. Mittel, 224
 - Besselsche, 198
 - Höldersche, 192
- Unterraum (s. a. Teilraum eines Vektor-
raums), 38
- Varianz, 24
 - generalisierte, 113
- Varianz-Kovarianz-Matrix
 - positiv semidefinit, 98
- Varianz-Kovarianzmatrix, 111
- Varianz-Kovarianzmatrizen
 - dyadisches Produkt, 75
- Vektoren
 - charakteristische, 95
- Vektorisierung, 165
- Vektornorm, 130, 132
- Vektorraum
 - n -dimensionaler, 47, 226
 - normierter, 131
 - Summe von V'räumen, 44
- Verbindungsvektor, 16
- vollständig, 193, 197
- vollständiges Funktionensystem, 199
- Vollständigkeitsrelation, 199
- vollstetig, 206
- Volterra-Typ, 207
- Zeilenrang, 78
- zentrierte Werte, 25
- Zentrierung, 118
- Zentrierungsmatrix, 75
- Zufallsvektor, 17