

Anwendung von p -Vektoren auf derivierte Matrizen.

Von B. GYIRES und O. VARGA in Debrecen.

Eine Matrix n -ter Ordnung $\mathfrak{A}$ kann als geordnetes Schema von n Vektoren a_i ($i, k=1, \dots, n$) eines n -dimensionalen Raumes aufgefaßt werden. Je p dieser Vektoren bestimmen einen p -Vektor. Durch Bildung der $N = \binom{n}{p}$ Kombinationen dieser Vektoren erhält man N p -Vektoren. Ordnet man diese p -Vektoren nach einer im Text (siehe (22)) angegebenen Art zu einem geordneten Schema an, so erhält man die p -te Derivierte von $\mathfrak{A}$, die kurz mit $C_p(\mathfrak{A})$ bezeichnet werden soll.

Ziel dieser Arbeit ist es nun zu zeigen, wie man mit Hilfe von p -Vektoren Sätze über derivierte Matrizen in geometrisch-anschaulicher Weise herleiten kann. Dieser Absicht entspricht es, daß wir nicht darauf abzielten, eine Reihe von neuen Sätzen über derivierte* Matrizen herzuleiten, sondern an Hand einiger Beispiele zeigen wollten, wie man umständliche elementare algebraische Betrachtungen durch geometrische Überlegungen ersetzen kann. Trotzdem scheinen Satz 1 und die Umkehrung in Satz 2 neu zu sein.

Im Folgenden setzen wir sämtliche auftretende Größen als reell voraus.

* * *

Ein p -dimensionales Parallellach — kurz p -Parallellach — des n -dimensionalen euklidischen Raumes E_n kann durch Angabe eines Punktes p_i und p linear unabhängige Vektoren a_i ($i=1, 2, \dots, p$) bestimmt werden.

Ergänzen wir die Vektoren a_i durch $n-p$ Einheitsvektoren $n_i, \dots, n_i$, die zu den a_i und außerdem paarweise untereinander orthogonal sind, zu einem n -Bein, so ist der Inhalt V_p des Parallellaches, durch den Betrag der Determinante

$$(1) \quad D_p = \begin{vmatrix} a_1 & \dots & a_p & n_1 & \dots & n_{n-p} \\ (1) & & (p) & (1) & & (n-p) \end{vmatrix}$$

erklärt, d. h. es ist

$$(2) \quad V_p = |D_p|.$$

Wegen der Voraussetzung über die Ergänzungsvektoren n_i , können wir das

Quadrat von V_p bei Verwendung der Cauchyschen Multiplikationsformel für rechteckige Matrizen in der Form

$$(3) \quad V_p^2 = \begin{vmatrix} a_i a_i & \dots & a_i a_i \\ (1)(1) & & (1)(p) \\ \cdot & \dots & \cdot \\ a_i a_i & \dots & a_i a_i \\ (p)(1) & & (p)(p) \end{vmatrix} = \begin{vmatrix} a_i a_i \\ (q)(q) \end{vmatrix} = \frac{1}{p!} P_{k_1 \dots k_p} P_{k_1 \dots k_p}$$

darstellen. Hier bedeutet $P_{k_1 \dots k_p}$ die durch

$$(4) \quad P_{k_1 \dots k_p} = \begin{vmatrix} a_{k_1} & \dots & a_{k_1} \\ (1) & & (p) \\ \cdot & \dots & \cdot \\ a_{k_p} & \dots & a_{k_p} \\ (1) & & (p) \end{vmatrix} = \begin{vmatrix} a_{k_1} & \dots & a_{k_p} \\ (q) & & (q) \end{vmatrix}$$

erklärte Determinante, und es ist über nicht eingeklammerte Zeiger, die zweimal vorkommen, von 1 bis n zu summieren.

Liegen zwei p -Parallelfäche in parallelen p -dimensionalen Ebenen E_p , so können wir für dieselben eine relative Orientierung einführen, indem wir festsetzen, daß bei Verwendung derselben Ergänzungsvektoren n_i ($i = 1, \dots, n-p$) die beiden p -Parallelfäche gleich orientiert sind, falls

$$(5) \quad \text{sign } D_p \text{ sign } D'_p = 1$$

ist, wenn D'_p die für das zweite p -Parallelfach gemäß (1) gebildete Determinante bedeutet. Sind b_i die das zweite p -Parallelfach aufspannenden Vektoren und wird

$$(6) \quad Q_{k_1 \dots k_p} = \begin{vmatrix} b_{k_1} & \dots & b_{k_p} \\ (q) & & (q) \end{vmatrix}$$

gesetzt, so ist die Bedingung (5) gleichwertig mit

$$(7) \quad \begin{vmatrix} a_i b_i \\ (q)(q) \end{vmatrix} = \frac{1}{p!} P_{k_1 \dots k_p} Q_{k_1 \dots k_p} > 0.$$

Unter einem einfachen p -Vektor — da nur solche vorkommen werden, lassen wir die Bezeichnung „einfach“ im Folgenden weg — verstehen wir die durch folgende drei Eigenschaften gekennzeichnete Gesamtheit von p -Parallelfächen¹⁾:

- 1) Zwei beliebige p -Parallelfäche liegen in parallelen Ebenen E_p .
- 2) Zwei beliebige p -Parallelfäche besitzen den gleichen Inhalt.
- 3) Zwei beliebige p -Parallelfäche sind gleich orientiert.

Bezeichnen wir ein beliebiges p -Parallelfach der Gesamtheit als Repräsentanten des p -Vektors, so folgt aus seiner Definition, daß dasselbe durch die zu einem Repräsentanten gehörigen $N = \binom{n}{p}$ Komponenten $P_{k_1 \dots k_p}(k_1, \dots, k_p)$

¹⁾ Wegen dieser Definition siehe etwa J. A. SCHOUTEN und D. J. STRUIK, Einführung in die neueren Methoden der Differentialgeometrie. I., (Groningen, 1935), S. 15–18.

durchlaufen sämtliche Kombinationen von 1 bis n) eindeutig bestimmt ist. Ist nämlich ein zweites p -Parallellach durch b_i bestimmt, so folgt wegen der Eigenschaft 1), daß

$$(8) \quad b_i = c_{\varrho\sigma} a_i, \quad |c_{\varrho\sigma}| \neq 0$$

ist. Daraus folgt wegen der Definition (4) und (8)

$$(9) \quad Q_{k_1 \dots k_p} = |c_{\varrho\sigma}| P_{k_1 \dots k_p}.$$

Aus den Eigenschaften 2) und 3) folgt daraus, daß

$$|c_{\varrho\sigma}| = 1$$

ist.

Wie für Vektoren im E_n , kann man für zwei p -Vektoren den Cosinus des von ihnen eingeschlossenen Winkels definieren. Dieser Winkel wird gleichzeitig den von zwei E_p eingeschlossenen Winkel bestimmen. Wir leiten die diesbezügliche bekannte Formel her²⁾, weil wir dazu Bemerkungen anschließen werden, die im Folgenden nötig sind. Die fragliche Definition ist eine einfache Verallgemeinerung der Definition des Cosinus des von zwei Strecken eingeschlossenen Winkels. Sind P und Q die beiden p -Parallellache, die den Punk p_i gemeinsam haben, dann projizieren wir das p -Parallellach P auf die Q enthaltende Ebene E_p . Ist $\bar{P}$ das projizierte p -Parallellach, $V_p(\bar{P})$ sein Inhalt, und $V_p(P)$ derjenige von P , dann soll der Cosinus des von P und Q eingeschlossenen Winkels φ durch

$$(10) \quad \cos \varphi = \frac{V_p(\bar{P})}{V_p(P)}$$

erklärt werden. Aus dem Beweis wird sich ergeben, daß diese Definition sinnvoll ist. Die Projektion $\bar{P}$ von P auf E_p ist dasjenige p -Parallellach, das von den Projektionen der Vektoren a_i auf E_p aufgespannt wird. Sind n_i ($i = 1, \dots, n-p$) zu E_p senkrechte Vektoren (die nicht Einheitsvektoren sein müssen), die b_i zu einen n -Bein von linear unabhängigen Vektoren ergänzen, so ergeben sich die projizierten Vektoren $\bar{a}_i$ aus der Zerlegung

$$(11) \quad a_i = \sum_{\sigma=1}^p c_{\varrho\sigma} b_i + \sum_{\tau=1}^{n-p} d_{\varrho\tau} n_i \quad (\varrho = 1, \dots, p)$$

in der Form

$$(12) \quad \bar{a}_i = \sum_{\sigma=1}^p c_{\varrho\sigma} b_i \quad (\varrho = 1, \dots, p).$$

²⁾ Wegen der im folgenden gegebenen Definition des Winkels zweier E_p vgl. B. SEGRE, Encyclopedie d. Math. Wiss. III., S. 801, für Bivektoren E. CARTAN, Lecons sur la géometrie des espaces, 2. Ausgabe (Paris, 1946) insb. S. 9.

Hieraus ergibt sich wegen der Inhaltsformel (3)

$$(13) \quad V_p^2(\bar{P}) = \frac{1}{p!} P_{k_1 \dots k_p} \bar{P}_{k_1 \dots k_p} = \frac{1}{p!} Q_{k_1 \dots k_p} Q_{k_1 \dots k_p} = |c_Q \sigma|^2 V_p^2(Q).$$

Multiplizieren wir (11) skalar mit b_i , so folgt hieraus

$$(14) \quad P_{k_1 \dots k_p} Q_{k_1 \dots k_p} = |c_Q \sigma| Q_{k_1 \dots k_p} Q_{k_1 \dots k_p}.$$

Aus (10) folgt nun wegen (13) und (14)

$$(15) \quad \cos \varphi = \frac{P_{k_1 \dots k_p} Q_{k_1 \dots k_p}}{\sqrt{P_{k_1 \dots k_p} P_{k_1 \dots k_p}} \sqrt{Q_{k_1 \dots k_p} Q_{k_1 \dots k_p}}} = \frac{\sum_{(k_1, \dots, k_p)} P_{k_1 \dots k_p} Q_{k_1 \dots k_p}}{\sqrt{\sum_{(k_1, \dots, k_p)} P_{k_1 \dots k_p}^2} \sqrt{\sum_{(k_1, \dots, k_p)} Q_{k_1 \dots k_p}^2}}.$$

Die eingeklammerten Zeiger unter dem Summenzeichen Σ bedeuten, daß $k_1, \dots, k_p$ sämtliche Kombinationen der Zahlen von 1 bis n durchlaufen. Durch Formel (15) ist nun auch der Cosinus der beiden p -Vektoren erklärt, für die P und Q je einen Repräsentanten bilden. Da (15) genau von der Form ist, wie die entsprechende Definition für Vektoren in einem N -dimensionalen Raum, so folgt auch hier, daß der absolute Betrag der rechten Seite von (15) höchstens gleich eins ist. Auf Grund der Formel (9) folgt, daß (15) unabhängig von der Wahl der p -Vektoren in ihrer E_p ist, sodaß (15) auch den Cosinus des Winkels der beiden E_p definiert. Der in Zähler von (15) auftretende Ausdruck

$$\sum_{(k_1, \dots, k_p)} P_{k_1 \dots k_p} Q_{k_1 \dots k_p}$$

soll als das skalare Produkt der beiden p -Vektoren bezeichnet werden. Untersuchen wir jetzt noch den Fall, für den dieses skalare Produkt verschwinden, also nach (15)

$$(16) \quad \cos \varphi = 0$$

ist. Wir behaupten, daß in diesem Falle in der linearen Vektormannigfaltigkeiten a_i mindestens eine eindimensionale Teilmannigfaltigkeit existiert, die zur Vektormannigfaltigkeit b_i senkrecht ist. Genauer beweisen wir folgenden

Hilfssatz: Sind unter den p -Vektoren $\bar{a}_i$, die die Projektionen von a_i auf die von den b_i bestimmte E_p bilden, $k \leq p$ linear unabhängig, dann gibt es in der Vektormannigfaltigkeit a_i eine $(p-k)$ -dimensionale Teilmannigfaltigkeit, die zu der von b_i orthogonal ist. Dann und nur dann verschwindet $\cos \varphi$, falls $k < p$ ist.

Beweis. Aus (12) folgt wegen der linearen Unabhängigkeit der b_i , daß es unter den Vektoren $\bar{a}_i$ dann und nur dann k linear unabhängige gibt,

falls der Rang von

$$(17) \quad \|c_{\varrho} \lambda\|$$

gleich k ist. Ist $k < p$, dann verschwindet $\cos \varphi$ wegen (10) und (13). Umgekehrt zieht das Verschwinden von $\cos \varphi$ wegen (14) und $V_p(Q) \neq 0$ die Ungleichung $k < p$ nach sich. Sind $\bar{a}_{(1)i}, \bar{a}_{(2)i}, \dots, \bar{a}_{(k)i}$ die k linear unabhängigen unter den Vektoren $\bar{a}_{(\varrho)i}$, dann gilt

$$(18) \quad \bar{a}_{(\varrho'')i} = \sum_{\varrho'=1}^k \lambda_{\varrho''\varrho'} \bar{a}_{(\varrho')i} \quad (\varrho'' = k+1, \dots, p).$$

Die von den $a_{(\varrho)i}$ bestimmte lineare Vektormannigfaltigkeit läßt sich auch durch die

$$(19) \quad \begin{aligned} a_{(\varrho')i}^* &= a_{(\varrho')i} & (\varrho' = 1, \dots, k) \\ a_{(\varrho'')i}^* &= \sum_{\varrho'=1}^k \lambda_{\varrho''\varrho'} a_{(\varrho')i} - a_{(\varrho'')i} & (\varrho'' = k+1, \dots, p) \end{aligned}$$

bestimmten Vektoren festlegen, da die Determinante der Vektortransformationen (19) von Null verschieden ist. Zerlegen wir nun die Vektoren $a_{(\varrho)i}^*$ nach dem n -Bein $b_{(\varrho)i}, n_{(\varrho)i}$ ($\varrho = 1, \dots, p; \tau = 1, \dots, n-p$), so folgt aus (11), (18) und (19)

$$(20) \quad \begin{aligned} a_{(\varrho')i}^* &= \sum_{\varrho=1}^p c_{\varrho'\varrho} b_{(\varrho)i} + \sum_{\tau=1}^{n-p} d_{\varrho'\tau} n_{(\tau)i} & (\varrho' = 1, \dots, k) \\ a_{(\varrho'')i}^* &= \sum_{\tau=1}^{k-p} \left(\sum_{\varrho'=1}^p \lambda_{\varrho''\varrho'} d_{\varrho'\tau} \right) n_{(\tau)i} & (\varrho'' = k+1, \dots, p). \end{aligned}$$

Die zweite Gruppe dieser Relationen besagt aber, daß die lineare Vektormannigfaltigkeit der $a_{(\varrho)i}^*$ eine $(p-k)$ -dimensionale Teilmannigfaltigkeit enthält, die zu derjenigen der $b_{(\varrho)i}$ orthogonal ist. Damit ist der Beweis der Hilfsatzes beendet.

Legen wir unter den $N = \binom{n}{p}$ Kombinationen von n Elementen p -ter Klasse eine willkürliche aber feste Reihenfolge fest, dann ist die p -te Derivierte $C_p(\mathfrak{A})$ der quadratischen Matrix n -ter Ordnung $\mathfrak{A}$, diejenige Matrix N -ter Ordnung, in deren i -ter Zeile und j -ter Spalte als Element diejenige Unterdeterminante p -ter Ordnung von $\mathfrak{A}$ auftritt, deren Zeilen und Spalten durch die i -te bzw. j -te Kombination der festgelegten Reihenfolge bestimmt ist.

Ordnet man der Matrix $\mathfrak{A}$ diejenigen n Vektoren zu, deren Komponenten durch die Spalten der Matrix gebildet werden, was durch die Schreibweise

$$(21) \quad \mathfrak{A} = \|a_{(1)i}, \dots, a_{(n)i}\|$$

angedeutet werden soll, dann wird die p -te Derivierte $C_p(\mathfrak{A})$ von $\mathfrak{A}$ in

folgender Weise mit Hilfe der zu je p dieser Vektoren gehörigen p -Vektoren darstellbar:

Ordnet man die p -Vektoren entsprechend der festgesetzten Reihenfolge der Kombinationen der a_i an, dann wird diejenige Matrix N -ter Ordnung, die man erhält, indem man jedem p -Vektor sein Komponentensystem $P_{k_1 \dots k_p}^{(S)}$ ($S=1, \dots, N$) in der Reihenfolge der Kombinationen der $k_1, \dots, k_p$ als Spalte zuordnet, gerade die p -te Derivierte $C_p(\mathfrak{A})$ von $\mathfrak{A}$.

Wir bringen diese Darstellung von $C_p(\mathfrak{A})$ durch die Schreibweise

$$(22) \quad C_p(\mathfrak{A}) = \left\| P_{k_1 \dots k_p}^{(1)}, P_{k_1 \dots k_p}^{(2)}, \dots, P_{k_1 \dots k_p}^{(N)} \right\|$$

zum Ausdruck.

Wir beweisen nun den

Satz 1. Die p -te Derivierte einer Matrix $\mathfrak{A}$ bestimmt dieselbe eindeutig, falls p ungerade ist, d. h. aus

$$(23) \quad C_p(\mathfrak{A}) = C_p(\mathfrak{B})$$

folgt $\mathfrak{A} = \mathfrak{B}$. Ist hingegen p gerade, so folgt aus (23) $\mathfrak{A} = \pm \mathfrak{B}$.

Beweis. Die Beziehung (23) bedeutet, daß in der Darstellung der beiden Matrizen in der Form (22) an gleicher Stelle stehende p -Vektoren gleich sind. Ist

$$\mathfrak{B} = \left\| b_i, \dots, b_i \right\|_{(1) \dots (n)}$$

so muß nach Eigenschaft 1) der p -Vektoren die von den $b_i, \dots, b_i$ aufgespannte lineare Vektormannigfaltigkeit mit der von den $a_i, \dots, a_i$ aufgespannten identisch sein, d. h. es muß jede der beiden Gruppen von Vektoren durch die andere linear kombinierbar sein. Halten wir a_i fest und bestimmen diejenigen $p-1$ linearen Vektormannigfaltigkeiten, bei denen sich die Basisvektoren von denjenigen von $a_i, a_i, \dots, a_i$ im 2-ten, 3-ten, bzw. p -ten Vektor unterscheiden und bestimmen dieselben Vektormannigfaltigkeiten mit den b_i , dann folgt bei Vergleichung von je zwei dieselbe Mannigfaltigkeit bestimmenden Systemen von Basisvektoren, daß

$$(24) \quad b_i = c_{k_1} a_i \quad (\text{nicht summieren über } k_1!)$$

sein muß. Da k_1 beliebig war, gilt dies somit für sämtliche Vektoren. Aus der Eigenschaft 2) der p -Vektoren folgt wegen (3) und (9), daß

$$(c_{k_1} c_{k_2} \dots c_{k_p})^2 = 1$$

für jede Kombination $k_1, \dots, k_p$ der Zahlen von 1 bis n ist, was mit

$$c_1^2 = c_2^2 = \dots = c_k^2 = 1$$

gleichbedeutend ist. Es kann daher bloß

$$(25) \quad b_{(r)} = \pm a_{(r)}$$

sein.

Nach der 3-ten Eigenschaft der p -Vektoren folgt aus Gleichung (1), daß die Gleichheit von zwei p -Vektoren erhalten bleibt, falls eine gerade Anzahl von Vektoren entgegengesetzt orientiert wird. Beachten wir die Bildung sämtlicher p -Vektoren aus den $b_{(k)}$, so folgt, daß für ungerades p auf der rechten Seite von (25) nur das positive Vorzeichen in Betracht kommt. Ist hingegen p gerade, dann kann

$$b_{(r)} = + a_{(r)}$$

und

$$b_{(r)} = - a_{(r)}$$

sein. Damit ist aber der Beweis von Satz 1 erbracht.

Als nächstes zeigen wir:

Mit $\mathfrak{A}$ ist auch $C_p(\mathfrak{A})$ regulär und umgekehrt.

Damit ist gleichbedeutend, daß mit $\mathfrak{A}^{-1}$ auch $C_p(\mathfrak{A})^{-1}$ existiert und umgekehrt.

Zum Beweis konstruieren wir $C_p(\mathfrak{A})^{-1}$. Die Existenz von $\mathfrak{A}^{-1}$ bedeutet, daß das zu $a_{(k)}$ adjungierte n -Bein $b_{(k)}$ existiert, für das

$$(26) \quad a_{(k)} b_{(s)} = \delta_{ks}$$

gilt. Diese beiden n -Beine bilden zwei zueinander polare n -Kanten, da z. B. die „Kante“ $b_{(1)}$ zu der aus den $a_{(2)}, \dots, a_{(n)}$ gebildeten $(n-1)$ -dimensionalen Seite des n -Kantes der $a_{(1)}$ senkrecht steht. Wir bilden nun aus den $b_{(k)}$ die N

p -Vektoren $Q_{i_1 \dots i_p}^{(s)}$ ($S=1, \dots, N$) und bestimmen mit ihnen, entsprechend der Darstellung (22), die p -te Derivierte der Matrix $\mathfrak{A}^{-1}$. Zwei p -Vektoren, die zu $C_p(\mathfrak{A})$ bzw. $C_p(\mathfrak{A}^{-1})$ gehören, aber in der angegebenen Reihenfolge nicht die gleiche Nummer besitzen, sind offensichtlich zueinander orthogonal, d. h. es gilt

$$(27) \quad \sum_{(i_1, \dots, i_k)} P_{i_1 \dots i_k}^{(S)} Q_{i_1 \dots i_k}^{(R)} = 0, R \neq S.$$

Ist hingegen $R=S$, dann wird wegen (26) und (27)

$$(28) \quad \sum_{(i_1, \dots, i_k)} P_{i_1 \dots i_k}^{(S)} Q_{i_1 \dots i_k}^{(S)} = 1.$$

(27) und (28) besagt gerade, daß

$$(29) \quad C_p(\mathfrak{A}^{-1}) = C_p(\mathfrak{A})^{-1}$$

ist. Existiert umgekehrt $C_p(\mathfrak{A})^{-1}$, dann muß

$$|C_p(\mathfrak{A})| \neq 0$$

sein. Daraus folgt aber, daß das n -Bein a_i der Matrix $\mathfrak{A}$ aus linear unabhängigen Vektoren besteht, da ja im entgegengesetzten Falle wegen (8) mindestens zwei Spalten proportional wären.

Als nächstes beweisen wir den

Satz 2. *Mit $\mathfrak{A}$ ist auch ihre p -te Derivierte orthogonal und umgekehrt.*

Der erste Teil dieses Satzes ist bekannt³⁾ und kann mit Hilfe unserer Überlegungen so gefolgert werden. Da die a_i ein normiertes orthogonales n -Bein bilden, so folgt aus der Inhaltsformel (3)

$$(30) \quad \sum_{(k_1, \dots, k_p)}^{(S)} P_{k_1 \dots k_p}^{(S)} P_{k_1 \dots k_p}^{(S)} = 1.$$

Ist hingegen $R \neq S$, dann folgt aus unserem Hilfssatz und der Orthogonalität der a_i

$$(31) \quad \sum_{(k_1, \dots, k_p)}^{(R)} P_{k_1 \dots k_p}^{(R)} P_{k_1 \dots k_p}^{(S)} = 0.$$

(30) und (31) sind die Orthogonalitätsbedingungen der Matrix (22). Setzen wir nun umgekehrt voraus, daß $C_p(\mathfrak{A})$ orthogonal ist, d. h. (30) und (31) gilt. Betrachten wir nun die beiden p -Vektoren, die durch

$$(\gamma) \quad a_i, a_i, \dots, a_i, \\ (1) \quad (2) \quad \quad \quad (\rho)$$

bzw.

$$(\delta) \quad a_i, \dots, a_i, a_i \quad (p+1 \leq \tau \leq n) \\ (1) \quad \quad \quad (\rho-1) \quad (\tau)$$

bestimmt sind. Da die beiden p -Vektoren orthogonal sind, und die linearen Vektormannigfaltigkeiten (γ) und (δ) $p-1$ linear unabhängige Vektoren gemeinsam haben, so ist die lineare Teilmannigfaltigkeit von (δ) , die gemäß des Hilfssatzes zu der von (γ) orthogonal ist, eindimensional, d. h. sie ist durch a_i

bestimmt. Die Vektoren der Gruppe (γ) und die Vektoren

$$(\varepsilon) \quad a_i, \dots, a_i \\ (\rho+1) \quad \quad \quad (\nu)$$

sind daher paarweise zueinander orthogonal. Durch entsprechende Wahl von zwei neuen p -Vektoren folgt aber, daß auch die Vektoren der Gruppe (γ) und (ε) untereinander paarweise orthogonal sind. Die a_i bilden also ein orthogonales n -Bein. Daß sie auch Einheitsvektoren sind, sieht man folgen-

³⁾ Siehe z. B. E. PASCAL: Repetitorium der höheren Mathematik. I. 1., 2. Aufl. (Berlin, 1910), S. 139, wo die erste Hälfte von unserem Satz 2 aus Lehrsätzen I und IV unmittelbar folgt.

dermaßen ein. Sind e_i Einheitsvektoren, für die

$$(32) \quad a_{(k)} = c_k e_{(k)} \quad (\text{nicht summieren über } k!)$$

ist, dann folgt aus (23), (9) und dem Umstand, daß die e_i Einheitsvektoren sind, daß

$$(33) \quad (c_{k_1} c_{k_2} \dots c_{k_p})^2 = 1$$

ist für jede Kombination p -ter Klasse $k_1, \dots, k_p$ der Zahlen von 1 bis n . Daraus folgt aber

$$c_k = \pm 1,$$

womit der Beweis von Satz 2 zu Ende geführt ist.

(Eingangen am 15. September 1951.)