

ADVANCED LINEAR ALGEBRA

DR. AMIT PRAKASH, DR. KAMLESH KUMAR,
MR. AJIT KUMAR



BlueRose ONE^{.com}
S t o r i e s M a t t e r
New Delhi • London

BLUEROSE PUBLISHERS

India | U.K.

Copyright © Dr. Amit Prakash, Dr. Kamlesh Kumar, Mr. Ajit Kumar 2026

All rights reserved by author. No part of this publication may be reproduced, stored in a retrieval system or transmitted in any form or by any means, electronic, mechanical, photocopying, recording or otherwise, without the prior permission of the author. Although every precaution has been taken to verify the accuracy of the information contained herein, the publisher assumes no responsibility for any errors or omissions. No liability is assumed for damages that may result from the use of information contained within.

BlueRose Publishers takes no responsibility for any damages, losses, or liabilities that may arise from the use or misuse of the information, products, or services provided in this publication.



For permissions requests or inquiries regarding this publication,
please contact:

BLUEROSE PUBLISHERS
www.BlueRoseONE.com
info@blurosepublishers.com
+91 8882 898 898
+4407342408967

ISBN: 978-93-7825-716-2

Cover design: Anuj
Typesetting: Tanya Raj Upadhyay

First Edition: May 2026

Preface

Linear Algebra stands as one of the most profound and far-reaching branches of modern mathematics, forming the intellectual backbone of diverse disciplines such as physics, computer science, engineering, economics, data science, and artificial intelligence. Its concepts not only provide powerful tools for solving practical problems but also reveal the inherent structure and symmetry underlying mathematical systems. Advanced Linear Algebra has been meticulously crafted with the aim of presenting this rich and elegant subject in a manner that is both rigorous and accessible, inspiring learners to develop a deep and lasting appreciation for its beauty and utility.

This book is the outcome of a sincere effort to bridge the gap between elementary understanding and advanced theoretical insight. It is designed primarily for undergraduate and postgraduate students of mathematics, as well as for aspirants preparing for competitive examinations. At the same time, it serves as a valuable reference for teachers and researchers seeking a structured and comprehensive exposition of the subject.

The text is systematically organized into nine chapters, each building upon the previous one to ensure a smooth and coherent progression of ideas. The journey begins with Matrices and Determinants, which establishes the algebraic foundation necessary for the study of linear systems. The next chapter, Linear Spaces, introduces vector spaces, subspaces, basis, and dimension—concepts that lie at the very heart of linear algebra.

The chapter on Systems of Linear Equations develops both theoretical and computational methods for analyzing and solving systems, highlighting their significance in real-world applications. Euclidean Spaces extends these ideas by incorporating inner products, norms, and orthogonality, thereby connecting algebra with geometry in a meaningful way.

In Linear Operators, the focus shifts to transformations and their properties, including eigenvalues, eigenvectors, and diagonalization—topics that play a crucial role in various applications. The subsequent chapter on Iterative Methods of Solving

Linear Systems and Eigenvalue Problems introduces numerical techniques that are indispensable in handling large-scale computational problems.

The study of Bilinear and Quadratic Forms provides deeper insight into the classification and transformation of forms, while the chapter on Tensors opens the gateway to multilinear algebra and its applications in advanced mathematics and theoretical physics. The final chapter, Elements of Group Theory, offers a concise yet meaningful introduction to algebraic structures, emphasizing the role of symmetry and abstraction in modern mathematical thought.

Throughout the book, emphasis has been placed on clarity of exposition, logical development of concepts, and mathematical rigor. Numerous examples have been included to illustrate abstract ideas, and a wide range of exercises—ranging from routine problems to more challenging ones—has been provided to enhance analytical skills and encourage independent thinking.

Every effort has been made to present the material in a student-friendly manner without compromising on depth and precision. The author firmly believes that mathematics is not merely a collection of formulas and techniques, but a disciplined way of thinking that cultivates clarity, creativity, and intellectual confidence.

It is sincerely hoped that this book will not only serve as a reliable academic resource but also ignite curiosity and enthusiasm among readers, motivating them to explore further the vast and fascinating landscape of linear algebra.

— Author

Table of Contents

Chapter – 1 Matrices and Determinants.....	1
Chapter – 2 Liner Spaces	41
Chapter – 3 Systems of Linear Equations	68
Chapter – 4 Euclidean Spaces	90
Chapter – 5 Linear Operators	120
Chapter – 6 Iterative Methods of Solving Linear Systems and Eigenvalue Problems.....	188
Chapter – 7 Bilinear and Quadratic Forms.....	221
Chapter – 8 Tensors	275
Chapter – 9 Elements of The Group Theory	319

CHAPTER – 1

MATRICES AND DETERMINANTS

In this chapter we consider tables of numbers which are called **matrices** and which play the most important role in what follows. The principal operations on matrices are introduced here and the properties of the so-called **determinants**, which are the main numerical characteristic of square matrices, are studied in detail.

1.1 Matrices

1.1.1 The concept of a matrix is a rectangular array of numbers containing a certain number m of rows and a certain number n of columns.

The numbers m and n are the **orders** of a matrix. When $m = n$, the matrix is called **square** and the number $m = n$ is its **order**.

In what follows, a matrix will be written either between double lines or parentheses:

$$\left\| \begin{array}{cccc} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \cdot & \cdot & \cdot & \cdot \\ a_{m1} & a_m & \dots & a_{mn} \end{array} \right\| \quad \text{or} \quad \left(\begin{array}{cccc} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \cdot & \cdot & \cdot & \cdot \\ a_{m1} & a_m & \dots & a_{mn} \end{array} \right)$$

For a brief designation of a matrix use will be often made either of a capital letter (say, A), or a symbol $\|a_{ij}\|$, sometimes followed by an elucidation:

$A = \|a_{ij}\| = (a_{ij})$ ($i = 1, 2, \dots, m; j = 1, 2, \dots, n$). The number a_{ij} the first index i is the number of the row and the second index j is the number of the column.

In the case of a square matrix

$$\left\| \begin{array}{cccc} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \cdot & \cdot & \cdot & \cdot \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{array} \right\| \quad (1.1)$$

the notions of a principal and a secondary diagonal are introduced. The **principal leading diagonal** of the matrix (1.1) is the diagonal $a_{11} a_{22} \dots a_{nn}$ passing from the upper left corner of the matrix to its lower right corner. The **secondary diagonal** of the same matrix is the diagonal $a_{n1} a_{(n-1)2} \dots a_{1n}$ passing from the lower left to the upper right corner.

1.1.2. The main operations on matrices and their properties. First of all, let us agree to consider two matrices to be **equal** if they have the same order and all their respective elements coincide.

We shall now define the main operations on matrices.

(a) **Addition of matrices.** The **sum** of two matrices $A = \|a_{ij}\| (i=1, 2, \dots, m; j=1, 2, \dots, n)$ and $B = \|b_{ij}\| (i=1, 2, \dots, m; j=1, 2, \dots, n)$ of the same orders m and n is a matrix $C = \|c_{ij}\| (i=1, 2, \dots, m; j=1, 2, \dots, n)$ of the same order m and n , whose elements c_{ij} are

$$c_{ij} = a_{ij} + b_{ij} (i=1, 2, \dots, m; j=1, 2, \dots, n). \quad (1.2)$$

The sum of two matrices is designated as $C = A + B$. The operation of forming the sum of matrices is called **addition**.

Thus, by definition,

$$\begin{aligned} & \left\| \begin{array}{cccc} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \cdot & \cdot & \cdot & \cdot \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{array} \right\| + \left\| \begin{array}{cccc} b_{11} & b_{12} & \dots & b_{1n} \\ b_{21} & b_{22} & \dots & b_{2n} \\ \cdot & \cdot & \cdot & \cdot \\ b_{m1} & b_{m2} & \dots & b_{mn} \end{array} \right\| \\ &= \left\| \begin{array}{cccc} (a_{11} + b_{11}) & (a_{12} + b_{12}) & \dots & (a_{1n} + b_{1n}) \\ (a_{21} + b_{21}) & (a_{22} + b_{22}) & \dots & (a_{2n} + b_{2n}) \\ \cdot & \cdot & \cdot & \cdot \\ (a_{m1} + b_{m1}) & (a_{m2} + b_{m2}) & \dots & (a_{mn} + b_{mn}) \end{array} \right\|. \end{aligned}$$

It immediately follows from the definition of the sum of matrices, to be more precise, from formula (1.2), that the operation of matrix addition has the same properties as the operation of addition of real numbers, namely:

- (1) commutativity : $A + B = B + A$.
- (2) associatively: $(A + B) + C = A + (B + C)$.

These properties free us from being troubled about the sequence of matrices when we add two or more matrices.

(b) **Multiplication of a matrix by a number.** The **product of the matrix** $A = \|a_{ij}\| (i=1, 2, \dots, m; j=1, 2, \dots, n)$ by a real number λ is a matrix $C = \|c_{ij}\| (i=1, 2, \dots, m; j=1, 2, \dots, n)$ whose elements c_{ij} are

$$c_{ij} = \lambda a_{ij} (i=1, 2, \dots, m; j=1, 2, \dots, n). \quad (1.3)$$

The notation $C = \lambda A$ or $C = A\lambda$ is used for the product of a matrix by a number. The operation of forming the product of a matrix by a number is called the **multiplication** of the matrix by that number.

It is clear from formula (1.3) that multiplication of a matrix by a number is subject to the following laws:

- (1) (1) associative law for the numerical factor: $(\lambda\mu)A = \lambda(\mu A)$;
- (2) distributive law for addition of matrices: $\lambda(A + B) = \lambda A + \lambda B$;
- (3) distributive law for the sum of the numbers: $(\lambda + \mu)A = \lambda A + \mu A$.

Remark. It is natural to term as the difference of two matrices A and B of the same order m and n matrix C of orders m and n which, being added to the matrix B , yields a matrix A . The natural notation for the difference of two matrices is $C = A - B$.

It is easy to verify that the difference C two matrices A and B can be obtained by the rule $C = A + (-1).B$.

(c) **Multiplication of matrices.** The **product** of the matrix $A = \|a_{ij}\| (i=1, 2, \dots, m; j=1, 2, \dots, n)$, of the respective orders m and n , by the

matrix $B = \|b_{ij}\| (i=1, 2, \dots, n; j=1, 2, \dots, p)$ of the respective orders n and p is a matrix $C = \|c_{ij}\| (i=1, 2, \dots, m; j=1, 2, \dots, p)$ of the respective orders m and p and i th elements c_{ij} defined by the formula

$$c_{ij} = \sum_{k=1}^n a_{ik} b_{kj} (i=1, 2, \dots, m; j=1, \dots, p). \quad (1.4)$$

The notation $C = A \cdot B$ is used to designate the product of the matrix A by the matrix B . The operation of forming the product of the matrix A by the matrix B is called the **multiplication** of those matrices.

It follows from the definition stated above that *the matrix A cannot be multiplied by any matrix B* . It is necessary that the number of columns of the matrix A be equal to the number of rows of the matrix B .

In particular, both products $A \cdot B$ and $B \cdot A$ can be defined only when the number of columns of A coincides with the number of rows of B and the number of rows of A coincides with the number of columns of B . In that case both matrices $A \cdot B$ and $B \cdot A$ are square but their orders, in a general case, are different. For the two products $A \cdot B$ and $B \cdot A$ to be not only defined but to have the same order, it is necessary and sufficient that both matrices A and B be square matrices of the same order.

Formula (1.4) is the rule of forming the elements of the matrix C which is the product of the matrix A by the matrix B . Here is the verbal formulation of this rule: *the element c_{ij} lying at the intersection of the i th row and the j th column of matrix $C = A \cdot B$ is equal to the sum of the pairwise products of the respective elements of the i th row of the matrix A and the j th column of the matrix B .*

We give the formula for multiplication of second-order square matrices as an example of application of the indicated rule:

$$\begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} \cdot \begin{vmatrix} b_{11} & b_{12} \\ a_{21} & b_{22} \end{vmatrix} = \begin{vmatrix} (a_{11}b_{11} + a_{12}b_{21}) & (a_{11}b_{12} + a_{12}b_{22}) \\ (a_{21}b_{11} + a_{22}b_{21}) & (a_{21}b_{12} + a_{22}b_{22}) \end{vmatrix}$$

Formula (1.4) yields the following properties of the product of the matrix A by the matrix B :

- (1) associative law: $(AB)C = A(BC)$;
 (2) distributive law for the sum of the matrices:

$$(A + B)C = AC + BC \text{ or } A(B + C) = AB + AC.$$

The distributive law follows immediately from formulas (1.4) and (1.2), and to prove the associative law it is sufficient to note that if

$$A = \|a_{ij}\| (i = 1, 2, \dots, m; j = 1, 2, \dots, n), \quad B = \|b_{ij}\| (i = 1, 2, \dots, n; j = 1, 2, \dots, p) \\ C = \|c_{ij}\| (i = 1, 2, \dots, p; j = 1, 2, \dots, r), \text{ then by virtue of (1.4) the element } d_{il} \text{ of the}$$

matrix $(AB)C$ is equal to $d_{il} = \sum_{R=1}^p \left(\sum_{j=1}^n a_{ij} b_{jk} \right) \cdot c_{Rl}$ and the element d_{il} of the

matrix $A(BC)$ is equal to $d_{il} = \sum_{j=1}^n a_{ij} \left(\sum_{R=1}^p b_{jk} c_{Rl} \right)$, but then the equality $d_{il} = d_{il}$

follows from the possibility of changing the order of summation with respect of j and R .

It is reasonable to pose the question concerning the commutative law for the product of the matrix A by the matrix B only for square matrices A and B of the same order (since, as was mentioned above, only for such matrices A and B are the products AB and BA defined and are the matrices of the same order). Simple examples show that *the product of two square matrices of the same order does not possess, in general, a commutative law*. Indeed, if we set

$$A = \begin{vmatrix} 0 & 1 \\ 0 & 0 \end{vmatrix}, B = \begin{vmatrix} 0 & 0 \\ 1 & 0 \end{vmatrix} \text{ then } AB = \begin{vmatrix} 1 & 0 \\ 0 & 0 \end{vmatrix}, \text{ and } BA = \begin{vmatrix} 0 & 0 \\ 0 & 1 \end{vmatrix}.$$

There are however, significant special cases in which the commutative law holds true*.

Among square matrices, there is a special class of **diagonal** matrices in each of which the elements lying outside the principal diagonal are zero. Every diagonal matrix of order n has the form

$$D = \begin{pmatrix} d_1 & 0 & \dots & 0 \\ 0 & d_2 & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & d_n \end{pmatrix},$$

where $d_1, d_2, \dots, d_n$ are arbitrary numbers. It is easy to see that when all these numbers are equal, i.e. $d_1 = d_2 = \dots = d_n = d$, then the equality $AD = DA$ is valid for any square matrix A of order n . Indeed, let us designated as c_{ij} and c'_{ij} the elements lying at the intersection of the i th row and the j th column of the matrices AD and DA respectively. Then, from equation (1.4) and the form of the matrix D we get

$$c_{ij} = a_{ij}d_j = a_{ij}d, \quad c'_{ij} = d_i a_{ij} = da_{ij}, \quad (1.6)$$

i.e. $c_{ij} = c'_{ij}$.

Among all the diagonal matrices (1.5) with coinciding elements $d_1 = d_2 = \dots = d_n = d$, an especially important role is played by the following two matrices. The first matrix is obtained at $d = 1$, is called an **identity** or **unit matrix** of order n and is denoted by I . The second matrix is obtained at $d = 0$, is called a null or **zero matrix** of order n and is denoted by O . Thus,

* Two matrices whose product the commutative law holds true are customarily called **commuting matrices**.

$$I = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & 1 \end{pmatrix}, \quad O = \begin{pmatrix} 0 & 0 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & 0 \end{pmatrix}.$$

By virtue of what was proved above, $AI = IA$ and $AO = OA$. Moreover, it is evident from (1.6) that

$$AI = IA = A, \quad AO = OA = O. \quad (1.7)$$

The first formula (1.7) characterizes the special role played by the identity matrix I , which is similar to that played by the number 1 in the multiplication of real numbers. As to the special role played by the null matrix O , it is revealed not only by the second formula (1.7) but also by the equality* $A + O = O + A = A$ which can be easily verified.

It should be pointed out in conclusion that the concept of a null matrix can be introduced for no square matrices as well (a null matrix is *any* matrix whose all elements are zero).

1.1.3. Hypermatrices. Suppose that by horizontal and vertical lines a matrix $A = \|a_{ij}\|$ is partitioned into separate rectangular cells each of which is a matrix of a smaller size and is called a **block** or a **submatrix** of the original matrix. In that case it is possible to regard the original matrix A as a certain new matrix $A = \|A_{\alpha\beta}\|$ (so-called **partitioned matrix** or **hypermatrix**) whose elements $A_{\alpha\beta}$ are the indicated blocks. Those elements are usually denoted by a capital letter to emphasize that they are usually denoted by a capital letter to emphasize that they are, in general, matrices and not numbers and (like ordinary numerical elements) are supplied with two indices, the first indicating the number of the row of the block and the second indicating the number of the column.

For example, the matrix

* This equality is a direct consequence of formula (1.2).

$$A = \left\| \begin{array}{cccc|cc} a_{11} & a_{12} & a_{13} & a_{14} & a_{15} & a_{16} \\ a_{21} & a_{22} & a_{23} & a_{24} & a_{25} & a_{26} \\ \hline a_{31} & a_{32} & a_{33} & a_{34} & a_{35} & a_{36} \\ a_{41} & a_{42} & a_{43} & a_{44} & a_{45} & a_{46} \\ a_{51} & a_{52} & a_{53} & a_{54} & a_{55} & a_{56} \end{array} \right\|$$

can be regarded as a hypermatrix $A = \left\| \begin{array}{cc} A_{11} & A_{12} \\ A_{21} & A_{22} \end{array} \right\|$, whose elements are the following blocks:

$$A_{11} = \begin{vmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \end{vmatrix}, A_{12} = \begin{vmatrix} a_{15} & a_{16} \\ a_{25} & a_{26} \end{vmatrix}$$

$$A_{21} = \begin{vmatrix} a_{31} & a_{32} & a_{33} & a_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \\ a_{51} & a_{52} & a_{53} & a_{55} \end{vmatrix}, A_{22} = \begin{vmatrix} a_{35} & a_{36} \\ a_{45} & a_{46} \\ a_{55} & a_{56} \end{vmatrix}.$$

It is remarkable that the principal operations on hypermatrices are performed according to the same rules that govern the operations on ordinary numerical matrices, with the only difference that blocks play the role of elements.

In fact, it can be easily verified that if the matrix $A = \|a_{ij}\|$ is a partitioned matrix with blocks $A_{\alpha\beta}$, then with the same partitioning, the matrix $\lambda A = \|\lambda a_{ij}\|$ is associated with the elements $\lambda A_{\alpha\beta}$.*.

*In that case the elements $\lambda A_{\alpha\beta}$ can be calculated by the rule of multiplication of the matrix $A_{\alpha\beta}$ by the number λ .

It can be as easily verified that if the matrices A and B are of the same order and are similarly partitioned, then the sum of the matrices A and B is associated with a hypermatrix with the elements $C_{\alpha\beta} + B_{\alpha\beta}$ (here $A_{\alpha\beta}$ and $B_{\alpha\beta}$ are the blocks of the matrices A and B).

Assume, finally, that A and B are two hypermatrices such that the number of columns in each block $A_{\alpha\beta}$ is equal to the number of rows of the block $B_{\beta\gamma}$ (so that product $A_{\alpha\beta}B_{\beta\gamma}$ is defined for any α, β and γ). Then the product $C = AB$ is a matrix with the elements $C_{\alpha\gamma}$, defined by the formula

$$C_{\alpha\gamma} = \sum_{\beta} A_{\alpha\beta} B_{\beta\gamma}.$$

To prove this formula, it is sufficient to write out its left-hand and right-hand sides in the terms of ordinary (numerical) elements of the matrices A and B (we leave this to the reader).

As an example of the application of hypermatrices, we shall consider the notion of the so-called **direct sum** of square matrices.

The **direct sum** of two square matrices A and B of orders m and n respectively is a square hypermatrix C of order $m+n$ equal to $C = \left\| \begin{array}{cc} A & O \\ O & B \end{array} \right\|$. The notation $C = A \oplus B$

is used to denote the direct summation and the operations of ordinary addition and multiplication of matrices can be easily verified with the aid of the properties of operations on hypermatrices:

$$(A_m \oplus A_n) + (B_m \oplus B_n) = (A_m + B_m) \oplus (A_n + B_n),$$

$$(A_m \oplus A_n) (B_m \oplus B_n) = A_m B_m \oplus A_n B_n$$

(in these formulas A_m and B_m are arbitrary square matrices of order m , and A_n and B_n are arbitrary square matrices of order n). We invite the reader to verify these formulas.

1.2 Determinants

The aim of this section is to construct the theory of determinants of any order n . We assume that the reader is acquainted with the second- and third-order determinants and shall make no references in our exposition. The acquaintance with the second- and third-order determinants will make it easier for the reader to comprehend the material that follows.

1.2.2 The concept of a determinant. Let us consider an arbitrary square matrix of any order n :

$$A = \left\| \begin{array}{cccc} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \cdot & \cdot & \cdot & \cdot \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{array} \right\|. \quad (1.8)$$

We shall associate with every matrix of this kind a definite numerical characteristic, called a determinant, corresponding to that matrix.

If the order n of the matrix (1.8) is equal to unity, then the matrix consists of one element a_{11} and the value of that element is the first-order determinant corresponding to that matrix.

If the order n of the matrix (1.8) is equal to two, if the matrix has the form

$$A = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix}, \quad (1.9)$$

then the second-order determinant corresponding to that matrix is a number equal to $a_{11}a_{22} - a_{12}a_{21}$ and designated as* $\Delta = \det A = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix}$. Thus, by definition,

$$\Delta = \det A = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} = a_{11}a_{22} - a_{12}a_{21}. \quad (1.10)$$

Formula (1.10) is a rule of forming a second-order determinant from the elements of the matrix corresponding to it. The verbal formulation of this rule is the following: the second-order determinant corresponding to matrix (1.9) is equal to the difference between the product of the elements on the principal diagonal the that matrix and the product of the elements lying on its secondary diagonal*.

In what follows, we shall speak of the elements, rows and columns of a determinant meaning the elements, rows and columns, respectively, of the matrix corresponding to that determinant.

Let us now elucidate the concept of a determinant of **any order** n , where $n \geq 2$. We shall introduce the notion of such a determinant by induction, assuming that we have already introduced the concept of a determinant of order $n - 1$, corresponding to an arbitrary square matrix of order $n - 1$.

*We shall agree to term as the **minor of any element** a_{ij} of the n th order matrix (1.8) the determinant of order $n - 1$ corresponding to the matrix obtained from matrix (1.8) by deleting the i th row and the j th column (the row and column at whose intersection the element a_{ij} is located).*

* As distinct from matrices, determinants are denoted by ordinary lines and not by double lines.

* Recall that the **principal** diagonal of a square matrix is the diagonal passing from the upper left corner to the lower right corner (i.e. $a_{11}a_{22}$ in the case of matrix (1.9)) and the secondary diagonal is that from the lower left to the upper right corner (i.e. $a_{11}a_{22}$ in the case of matrix (1.9)).

The minor of the element a_{ij} is designated as $\bar{M}_j^i$. The superscript denotes the number of the row, the subscript denotes the number of the column and the bar over the letter M denotes that the indicated row and column are deleted.

A **determinant of order n** corresponding to matrix (1.8) is a number equal to

$\sum_{j=1}^n (-1)^{1+j} a_{1j} \bar{M}_j^1$ and designated as

$$\Delta = \det A = \begin{vmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{vmatrix}. \quad (1.11)$$

Thus, by definition

$$\Delta = \det A = \begin{vmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{vmatrix} = \sum_{j=1}^n (-1)^{1+j} a_{1j} \bar{M}_j^1. \quad (1.12)$$

Formula (1.12) is the rule of forming a determinant of order n from the elements of the first row of the corresponding matrix and from the minors $\bar{M}_j^1$ of the elements of the first row which are determinants of order $n - 1$.

Note that for $n = 2$ rule (1.12) exactly coincides with rule (1.10) since in that case the minors of the elements of the first row have the form $\bar{M}_1^1 = a_{22}, \bar{M}_2^1 = a_{21}$.

The natural question arises whether it is possible to use the elements and the corresponding minors of an arbitrary i th row matrix (1.8) rather than of the first row in order to obtain the value of determinant (1.11). The following *fundamental* theorem provides an answer to that question.

Theorem 1.1 *Whatever the number of the row i ($i = 1, 2, \dots, n$) the following formula* is valid for the n th-order determinant (1.11) :*

$$\Delta = \det A = \sum_{j=1}^n (-1)^{i+j} a_{ij} \bar{M}_j^i, \quad (1.13)$$

called the expansion of that determinant according to the i th row.

Remark. It should be emphasized that in this formula the degree to which the number (-1) is raised is equal to the sum of the numbers of the row and the column at whose intersection the element is located.

The proof of theorem 1.1. Formula (1.13) should be proved only for the numbers $i = 2, 3, \dots, n^*$. For $n = 2$ (i.e. for a second-order determinant) this formula should be proved only for the number $i = 2$, i.e. for $n = 2$ it is only necessary to prove the formula

$$\Delta = \det A = \sum_{j=1}^2 (-1)^{2+j} a_{2j} \bar{M}_j^2, = -a_{21} \bar{M}_1^2 + a_{22} \bar{M}_2^2.$$

The validity of the last formula follows immediately from the expressions for the minors of matrix (1.9) $\bar{M}_1^2 = a_{12}, \bar{M}_2^2 = a_{11}$, by virtue of which the right-hand side of the formula coincides with the right-hand side of (1.10). We have thus proved the theorem for $n = 2$.

We shall vary out the proof of formula (1.13) for an arbitrary $n > 2$ by induction, i.e. we shall assume that the formula of the form (1.13) of the expansion according to any row is valid for a determinant of order $n-1$, and proceeding from this assumption we shall make sure that formula (1.13) is also valid for a determinant of order n .

* By the meaning of the theorem $n \geq 2$.

To prove this, we shall need the notion of minors of matrix (1.8) of order $n-2$. The determinant of order $n-2$ corresponding to the matrix which can be obtained from matrix (1.8) by deleting two rows with the numbers i_1 and i_2 and two columns with the

numbers j_1 and j_2 is called a **minor of order** $n-2$ and is designated as $\bar{M}_{j_1 j_2}^{i_1 i_2}$.

An n th-order determinant Δ is introduced by formula (1.12), and in this formula every minor $\bar{M}_j^1$ is a determinant of order $n-1$ for which the formula the form (1.13) of the expansion according to any row is valid by assumption.

Having fixed any number $i (i = 2, 3, \dots, n)$, we expand every minor $\bar{M}_j^1$ in formula (1.12) *according to the i th row* of the principal determinant (1.11) (in the minor $\bar{M}_j^1$ itself that row is the $(i-1)$ th).

As a result the whole determinate Δ will be represented as a certain linear combination (*Recall that a linear combination of some quantities is the sum of the products of those quantities by certain real numbers) of $(n-2)$ th-order minors $\bar{M}_{ik}^{li}$ with the noncoinciding number j and R , i. e. in the form

$$\Delta = \sum_{j=1}^n \sum_{k=1} \theta_{jk} \bar{M}_{ik}^{li}. \quad (1.14)$$

To calculate θ_{ik} we note that the minor $\bar{M}_{jk}^{li}$ results from the expansion according to the $(i-1)$ th row of *only the following two minors of order $n-1$ that correspond to the elements of the first row of matrix (1.8):* the minor $\bar{M}_j^1$ and the minor $\bar{M}_k^1$ (since only these two minors of the elements of the first row contain **all the columns** of the minor $\bar{M}_k^1$).

In the expansion of the minors $\bar{M}_j^1$ and $\bar{M}_k^1$ according to the indicated $(i-1)$ th row, we write out only those summands which contain the minor $\bar{M}_{jk}^{li}$ (and denote by

dots the other summands which are of no interest to us). Bearing in mind that the element a_{iR} of the minor $\bar{M}_j^1$ is at the intersection of the $(i-1)$ th row and the $(R-1)$ th column of the minor, and the element a_{ij} of the minor $\bar{M}_R^1$ is at the intersection of the $(i-1)$ th row and the j th column of that minor, we get

$$\bar{M}_k^1 = (-1)^{(i-1)+(k-1)} a_{ik} \bar{M}_{jk}^{1i} + \dots, \quad (1.15)$$

$$\bar{M}_r^1 = (-1)^{(i-1)+j} a_{ik} \bar{M}_{jk}^{1i} + \dots \quad (1.16)$$

Substituting (1.15) and (1.16) into the right-hand side of (1.12) and collecting the coefficients in $\bar{M}_{jk}^{1i}$, we find that the multiplier θ_{jk} in (1.14) has the form

$$\theta_{jk} = (-1)^{(1+i+j+k)} [a_{1j} a_{ik} - a_{1k} a_{ij}]. \quad (1.17)$$

To complete the proof of the theorem, we shall show that the right-hand side of (1.13) is also equal to the sum appearing on the right-hand side of (1.14) with the same values of (1.17) for θ_{jk} .

For that purpose, we expand every minor $\bar{M}_j^i$ of order $n-1$ on the right-hand side of (1.13) **according to the first row**. As a result, the whole-hand side of (1.13) will become a linear combination of the same minors $\bar{M}_{jk}^{1i}$

$$\sum_{j=1}^n \sum_{k < j} \theta_{jk} \bar{M}_{jk}^{1i} \quad (1.18)$$

With some coefficients θ_{jk} , and we have now to calculate the multipliers θ_{jk} , and ascertain the validity of formula (1.17) for them.

To do that, we note that minor $\bar{M}_{jk}^{1i}$ result from the expansion according to the first row of only the following two minors of matrix (1.8): the minor $\bar{M}_{jk}^{1i}$ and the minor

$\bar{M}_k^i$ (Since only these two minors of the elements of the i th row contain **all the columns** of the minor $\bar{M}_{jk}^{li}$).

In the expansions of the minors $\bar{M}_j^i$ and $\bar{M}_k^i$ according to the first row, we write out only those summands which contain the minor $\bar{M}_{jk}^{li}$ (and denote by dots the other summands which are of no interest to us). Taking into account that the element a_{1k} of the minor $\bar{M}_j^i$ lies at the intersection of the first row and the $(k-1)$ th column* of that minor and the element a_{1j} of the minor $\bar{M}_k^i$ lies at the intersection of the first row and the j th column** of that minor, we get

$$\bar{M}_j^i = (-1)^{1+(k-1)} a_{1j} \bar{M}_{jk}^{li} + \dots, \quad (1.19)$$

$$\bar{M}_k^i = (-1)^{1+j} a_{1k} \bar{M}_{jk}^{li} + \dots \quad (1.20)$$

Substituting (1.19) and (1.20) into the right-hand side of (1.13) and collecting the coefficients in $\bar{M}_{jk}^{li}$, we find that θ_{ik} in the sum (1.18) is defined by the same formula (1.17) as in equation (1.14).

We have proved Theorem 1.1.

*Since $j < k$ and the j th column of matrix (1.8) is absent in the minor $\bar{M}_k^i$.

**Since $j < k$ and only the k th column of matrix (1.8) is absent in the minor $\bar{M}_j^i$.

Theorem 1.1 has established the possibility of expanding an n th-order determinant according to any of its rows. A question now arises whether it is possible to expand a determinant of order n according to any of its **columns**. The following *fundamental* theorem provides a positive answer.

Theorem 1.2. *Whatever the number of the column j ($j=1,2,\dots,n$), the following formula holds true for the n th-order determinant (1.11):*

$$\Delta = \det A = \sum_{i=1}^n (-1)^{i+1} a_{ij} \bar{M}_j^i. \quad (1.21)$$

This formula is known as the **expansion of the determinant according to the j th column**.

Proof. It is sufficient to prove the theorem for $j=1$, i.e. to derive an expansion formula according to the first column

$$\Delta = \sum_{i=1}^n (-1)^{i+1} a_{i1} \bar{M}_1^i. \quad (1.22)$$

When formula (1.22) is derived, then to prove formulas (1.21) for any $j=2, 3, \dots, n$, it is sufficient to interchange the roles of the rows and columns and to repeat word for word the arguments of Theorem 1.1.

Formula (1.22) can be derived by induction.

It is easy to verify this formula for $n=2$ (since for $n=2$ the minors of the elements of the first column have the form $\bar{M}_1^1 = a_{22}, \bar{M}_1^2 = a_{12}$, it follows that for $n=2$ the right-hand side of (1.22) coincides with the right-hand side of (1.10)).

Let us assume that the formula for the expansion according to the first column (1.22) is valid for a determinant of order $n-1$ and, proceeding from this fact, ascertain the validity of that formula for a determinant of order n .

With this aim in view, we isolate the first term $a_{11} \bar{M}_1^1$ on the right-hand side of formula (1.12) for a determinant Δ of order n and expand the $(n-1)$ th-order determinant $\bar{M}_j^1$ **according to the first column** in all the other terms.

As a result, formula (1.12) assumes the form

$$\Delta = a_{11} \bar{M}_1^1 + \sum_{j=2}^n \sum_{i=1}^n \theta_{ij} \bar{M}_{ij}^{1i}, \quad (1.23)$$

where θ_{ij} are certain coefficients which have to be defined. To calculate θ_{ij} , we note that the minor $\bar{M}_{ij}^{1i}$ results from the expansion according to the first

column of only one of the minors of order $n - 1$ corresponding to the first row*, of the minor $\bar{M}_j^1$. In the expansion of the minor $\bar{M}_i^1$ (for $j \geq 2$) according to the first column we write only that term which contains the minor $\bar{M}_{1j}^{1i}$ (and use dots to denote the other terms which are of no interest to us). Taking into account that the element a_{ij} of the minor $\bar{M}_j^1$ (for $\bar{M}_j^1$) lies in the $(i - 1)$ th row and the first column of that minor, we find that for $j \geq 2$

$$\bar{M}_j^1 = (-1)^{(i-1)+1} a_{il} \bar{M}_{1j}^{1i} + \dots \quad (1.24)$$

Substituting (1.24) into the right-hand side of (1.12) (from which the first term has been deleted) and collecting the coefficients in $\bar{M}_{1j}^{1i}$, we find that the coefficient θ_{ij} in (1.23) has the form

$$\theta_{ij} = (-1)^{i+j+1} a_{il} a_{il}. \quad (1.25)$$

It remains to prove that the right-hand side of (1.22) is also equal to the sum appearing on the right-hand side of (1.23) with the same value (1.25) for θ_{ij} .

For that of (1.22) and expand the $(n - 1)$ th-order minor $\bar{M}_1^i$ **according to the first row** in all other term.

* The minor $\bar{M}_1^1$ is assumed to be excluded.

As a result, the right-hand side of (1.22) will be the sum of the first term $a_{11} \bar{M}_1^1$ and a linear combination with some coefficients θ_{ij} of the minors $\bar{M}_{1j}^{1i}$ of order $n - 2$, i.e. will have the form

$$a_{11} \bar{M}_1^1 + \sum_{l=2}^n \sum_{j=2}^n \theta_{ij} \bar{M}_{1j}^{1i}, \quad (1.26)$$

and we have only to calculate the terms θ_{ij} and ascertain the expansion, according to the first row, of only one of the $(n - 1)$ th-order minors corresponding to the first

column, the minor $\bar{M}_1^i$. In the expansion of the minor $\bar{M}_1^i$ (for $i \geq 2$) according to the first row we write only that term which contains the minor $\bar{M}_{1j}^{li}$ (and denote by dots the other terms which are of no interest to us). Taking into account that the element a_{ij} of the minor $\bar{M}_1^i$ is in the first row and the $(j-1)$ th column of that minor, we find that for $i \geq 2$

$$\bar{M}_1^i = (-1)^{1+(-1)} a_{1j} \bar{M}_{1j}^{li} + \dots \quad (1.27)$$

Substituting (1.25) into the right-hand side of (1.22) from which the first term has been deleted, and forming a coefficient in $\bar{M}_{1j}^{li}$, by collecting the terms we find that θ_{ij} in the sum (1.26) is defined by the same formula (1.25) as in equation (1.23). We have proved Theorem 1.2.

1.2.2 Expressing a determinant directly in terms of its elements. Let us derive a formula making it possible to express a determinant of order n directly in terms of its elements (by-passing the minors).

Suppose each numbers $\alpha_1, \alpha_2, \dots, \alpha_n$ assumes one of the values $1, 2, \dots, n$, none of these numbers coinciding (in that case we say that the numbers re a certain **transposition** of the numbers $\alpha_1, \alpha_2, \dots, \alpha_n$ are a certain **transposition** of the numbers $1, 2, \dots, n$). We form all possible pairs $\alpha_i \alpha_j$ of the numbers $\alpha_1, \alpha_2, \dots, \alpha_n$ and say that the pair $\alpha_i \alpha_j$ forms a **disorder** if $\alpha_i > \alpha_j$ for $i < j$. The total number of disorders formed by these pairs which can be compiled from the numbers $\alpha_1, \alpha_2, \dots, \alpha_n$.

Next we derive by induction the following formula for the n th-order determinant (1):

$$\Delta = \det A = \sum_{\alpha_1, \alpha_2, \dots, \alpha_n} (-1)^{N(\alpha_1, \alpha_2, \dots, \alpha_n)} a_{\alpha_1 1} a_{\alpha_2 2} \dots a_{\alpha_n n} \quad (1.28)$$

(in this formula the summation extends over all possible interchanges (transpositions) $\alpha_1, \alpha_2, \dots, \alpha_n$ of the numbers $1, 2, \dots, n$; the number of these transpositions is evidently $n!$).

Formula (1.28) can be easily verified when $n = 2$ (then only two transpositions, 1, 2 and 2, 1 are possible and, since $N(1, 2) = 0$, $N(2, 1) = 1$, formula (1.28) turns into equation (1.10)).

To prove the formula by induction, we assume that at $n > 2$ formula (1.28) holds true for a determinant of order $n - 1$.

Then, writing the expansion of the n th-order determinant (1.11) according to the first column*

$$\Delta = \det A = \sum_{\alpha_1=1}^n (-1)^{\alpha_1+1} \alpha_1 \overline{M}_1^{\alpha_1}, \quad (1.29)$$

we can, by virtue of the induction hypothesis, represent each minor $\overline{M}_1^{\alpha_1}$ order $n - 1$ in the form

* This time it is convenient for us to designate the index of summation as α_1 .

$$\overline{M}_1^{\alpha_1} = \sum_{\alpha_2, \dots, \alpha_n} (-1)^{N(\alpha_2, \dots, \alpha_n)} a_{\alpha_2 2} \dots a_{\alpha_n n} \quad (1.30)$$

(the summation extends over all possible transpositions $\alpha_2, \dots, \alpha_n$ of $(n - 1)$ natural numbers from 1 to n with the exception of the number α_1).

Since besides the pairs formed from the numbers $\alpha_2, \dots, \alpha_n$ we can form only the pairs $\alpha_1, \alpha_2, \dots, \alpha_n$ and since among the numbers α_1, α_n from the numbers $\alpha_1, \alpha_2, \dots, \alpha_n$ and since among the numbers $\alpha_2, \dots, \alpha_n$ there are exactly $(\alpha_1 - 1)$ numbers smaller than the number α_1 , it follows that $N(\alpha_1, \alpha_2, \dots, \alpha_n) = N(\alpha_2, \dots, \alpha_n) + \alpha_1 - 1$.

Hence $(-1)^{N(\alpha_2, \dots, \alpha_n)} (-1)^{\alpha_1+1} = (-1)^{N(\alpha_1, \alpha_2, \dots, \alpha_n)}$ and, substituting (1.30) into (1.29), we exactly obtain formula (1.28). We have thus completed the derivation of formula (1.28).

It should be pointed out in conclusion that in the majority of courses of linear algebra out in conclusion that in the majority of courses of linear algebra formula (1.28) underlies the notion of an n th-order determinant.

1.2.3. Laplace's expansion. We shall derive here a remarkable formula generalizing the formula for expansion of a determinant of order n according to some row.

For that purpose we introduce into consideration **two types** of minors of a matrix of order n (1.8).

Suppose k is any number smaller than n and $i_1, i_2, \dots, i_k$ and $j_1, j_2, \dots, j_k$ are arbitrary numbers satisfying the conditions $1 \leq i_1 < i_2 < \dots < i_k \leq n, 1 \leq j_1 < j_2 < \dots < j_k \leq n$.

Minors of the **first** type $M_{j_1 j_2 \dots j_k}^{i_1 i_2 \dots i_k}$ are determinants of order k corresponding to the matrix formed by the elements of matrix (1.8) lying at the intersection of k rows labelled $i_1, i_2, \dots, i_k$ and k columns labelled $j_1, j_2, \dots, j_k$.

Minors of the **second** type $\overline{M}_{j_1 j_2 \dots j_k}^{i_1 i_2 \dots i_k}$ are determinants of order $n - k$ corresponding to the matrix obtained from matrix (1.8) by deleting R rows labelled $i_1, i_2, \dots, i_k$ and k columns labelled $j_1, j_2, \dots, j_k$.

Minors of the **second** types are naturally called **complementary** with respect to minors of the first types.

Theorem 1.3 (Laplace's expansion). *For any number R smaller than n and for any fixed numbers of rows $i_1, i_2, \dots, i_k$ such that $1 \leq i_1 < i_2 < \dots < i_k \leq n$, the following formula holds true for the determinant of order n (1.11):*

$$\Delta \det A = \sum_{j_1, j_2, \dots, j_k} (-1)^{i_1 + \dots + j_k + j_1 + \dots + j_R} \times M_{j_1 j_2 \dots j_k}^{i_1 i_2 \dots i_k} \overline{M}_{j_1 j_2 \dots j_k}^{i_1 i_2 \dots i_k}. \quad (1.31)$$

This formula is called the **expansion of the determinant according to k rows** $1 \leq j_1 < j_2 < \dots < j_k \leq n$.

Proof. Note first of all that formula (1.31) is a generalization of the formula we have proved for the expansion of a determinant for order n according to **one** of its rows labelled i_1 into which it passes at $k=1$ (in that case the minor $M_{j_1}^{i_1}$ coincides with the element $a_{i_1 j_1}$ and minor $\overline{M}_{j_1}^{i_1}$ is the minor of the element $a_{i_1 j_1}$ introduced above).

Thus we have proved formula (1.31) for $k=1$. We shall use induction to prove this formula for any k satisfying the inequalities $1 < k < n$, i.e. we shall assume that formula (1.31) is valid for $k-1$ rows and proceeding from this fact ascertain the validity of formula (1.31) for k rows.

Thus suppose that $1 < k < n$ and that any k rows of matrix (1.8) labelled $i_1, i_2, \dots, i_k$ and satisfying the condition $1 \leq i_1 < i_2 < \dots < i_k \leq n$ are fixed. Then, by assumption, the following formula holds true for $k-1$ rows labelled $i_1, \dots, i_{k-1}$:

$$\Delta = \sum_{j_1, \dots, j_{k-1}} (-1)^{i_1 + \dots + i_{k-1} + j_1 + \dots + j_{k-1}} \times M_{j_1 \dots j_{k-1}}^{i_1 \dots i_{k-1}} \overline{M}_{j_1 \dots j_{k-1}}^{i_1 \dots i_{k-1}}. \quad (1.32)$$

(the summation here is taken over all possible values of the indices $j_1, \dots, j_{k-1}$ satisfying the conditions $1 \leq j_1 < j_2 < \dots < j_{k-1} \leq n$).

Let us expand every minor $\overline{M}_{j_1 \dots j_{k-1}}^{i_1 \dots i_{k-1}}$ in formula (1.32) according to the row labelled i_k in matrix (1.8). As a result, the whole determinant Δ will be represented as some linear combination of the minors $\overline{M}_{j_1 \dots j_{k-1} j_k}^{i_1 \dots i_{k-1} i_k}$ with coefficients which we shall designate as $\theta_{j_1 \dots j_k}$, i.e. the equality* $\Delta = \sum_{j_1, \dots, j_k} \theta_{j_1 \dots j_k} \overline{M}_{j_1 \dots j_k}^{i_1 \dots i_k}$ will hold

true for Δ and it remains to calculate the coefficients $\theta_{j_1 \dots j_k}$ and ascertain that they are equal to

$$\theta_{j_1 \dots j_k} = (-1)^{i_1 + \dots + i_k + j_1 + \dots + j_k} M_{j_1 \dots j_k}^{i_1 \dots i_k} \quad (1.33)$$

With that aim in view, we note that the $(n-k)$ th minor $\overline{M}_{j_1 \dots j_k}^{i_1 \dots i_k}$

* As before, the summation in this equation is taken over all possible values of the indices $j_1, \dots, j_k$ satisfying the condition $1 \leq j_1 < j_2 < \dots < j_k \leq n$.

results from the expansion, according to the row labelled i_k , of only the following k minors of the $(n-k+1)$ th order:

$$\overline{M}_{j_1 \dots j_k \text{ (without } j_s)^{**}}^{i_1 \dots i_{k-1}} \quad (s = 1, 2, \dots, k), \quad (1.34)$$

since none of the other minors of order $n-k+1$ containing the row i_s , contains all the rows and all the columns of the minor $\overline{M}_{j_1 \dots j_k}^{i_1 \dots i_k}$.

In the expansion of every minor (1.34) according to the row labelled i_k of matrix (1.8), we write out only the term containing the minor $\overline{M}_{j_1 \dots j_k}^{i_1 \dots i_k}$ (and denote by dots all the other terms which are of no interest to us). Taking into account that in every minor (1.34) the element a_{ikj_s} lies in the row $[i_k - (k-1)]$ and the column $[j_s - (s-1)]$ of that minor, (The designation $\overline{M}_{j_1 \dots j_k \text{ (without } j_s)}^{i_1 \dots i_{k-1}}$ is used for the minor corresponding to the intersection of the rows labelled $i_1, \dots, i_{k-1}$ and the columns labelled $j_1, j_2, \dots, j_k$, except for the column labelled j_s , and the designation (1.34) is used for the complementary minor.)

we get

$$\overline{M}_{j_1 \dots j_k \text{ (without } j_s)}^{i_1 \dots i_{k-1}} = (-1)^{[i_k - (k-1)] + [j_s - (s-1)]} a_{ikj_s} \overline{M}_{j_1 \dots j_k}^{i_1 \dots i_k} + \dots$$

We have now only to bear in mind that every minor (1.34) in formula (1.32) must be multiplied by the factor

$$(-1)^{(i_1 + \dots + i_{k-1}) + (j_1 + \dots + j_k) - j_s} M_{j_1 \dots j_k \text{ (without } j_s)}^{i_1 \dots i_{k-1}}$$

and then the sum is taken over all s taken over all s from 1 to k . Taking also into consideration that $(-1)^{2-2k-2s} = 1$, we find that

$$\theta_{j_1 \dots j_k} = (-1)^{i_1 + \dots + i_k + j_1 + \dots + j_k} \left[\sum_{s=1}^k (-1)^{k+s} M_{j_1 \dots j_k \text{ (without } j_s)}^{i_1 \dots i_{k-1}} \right].$$

Nothing that the sum in brackets is the expansion of the minor $M_{j_1 \dots j_k}^{i_1 \dots i_k}$ according to its last, k th, row, we finally get formula (1.33) for $\theta_{j_1 \dots j_k}$. We have completed the proof of Laplace's theorem.

Remark. The formula for the expansion of a determinant according to some k columns is written and derived by a complete analogy with formula (1.32).

1.2.4. Properties of determinants. Given below are some properties possessed by an arbitrary n th-order determinant.

1°. The property of the equivalence of rows and columns. The **transposition** of any matrix or determinant is an operation as a result of which the rows and columns are interchanged but their sequence is retained. The transposition of a matrix A results in a matrix **transposed with respect to the matrix A** and designated as A' .

In what follows we agree to use the symbols $|A|, |B|, |A'| \dots$ for the determinants of the square matrices $A, B, A' \dots$ respectively. The first property of a determinant is formulated as follows: *in a transposition the value of a determinant is retained, i.e. $|A'| = |A|$.*

This property follows directly from Theorem 1.2 (it is sufficient to note that the expansion of the determinant $|A|$ according to the first column identically coincides with that of the determinant $|A'|$ according to the first row).

The property we have proved signifies a complete equivalence of the rows and columns and enables us to establish all the other properties **only for the rows** and to be sure of their validity for the columns.

2°. The property of antisymmetry in interchanging two rows (or two columns). *When two rows (or two columns) are interchanged, the determinant retains its absolute value but changes its sign to the opposite.*

For a second-order determinant, this property can be easily verified (it follows immediately from (1.10) that the determinants $\begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix}$ and $\begin{vmatrix} a_{21} & a_{22} \\ a_{11} & a_{12} \end{vmatrix}$ differ only in sign).

Assuming that $n > 2$, we shall consider the n th-order determinant (1.11) and suppose that two rows labelled i_1 and i_2 are interchanged in that determinant. Writing Laplace's formula for expansion according to those two rows, we have

$$\Delta = \sum_{j_1, j_2} (-1)^{i_1+i_2+j_1+j_2} M_{j_1 j_2}^{i_1 i_2} \overline{M}_{j_1 j_2}^{i_1 i_2}. \quad (1.35)$$

When the rows labelled i_1 and i_2 are interchanged, every second-order determinant $M_{j_1 j_2}^{i_1 i_2}$ changes sign to the opposite by virtue of what was proved above, and all the other quantities appearing under the summation sign in (1.35) do not depend on the elements of the rows labelled i_1 and i_2 and retain their values. We have thus proved property 2°.

3°. Linear property of a determinant. *We say that some row $(a_1, a_2, \dots, a_n)$ is a linear combination of the rows $(b_1, b_2, \dots, b_n)$, $(c_1, c_2, \dots, c_n), \dots, (d_1, d_2, \dots, d_n)$ with coefficients $\lambda, \mu, \dots, \nu$ if $a_j = \lambda b_j + \mu c_j + \dots + \nu d_j$ for all $j = 1, 2, \dots, n$.*

The linear property of a determinant can be formulated as follows: *if in the n th-order determinant Δ a certain i th row $(a_{i1}, a_{i2}, \dots, a_{in})$ is a linear combination of two rows $(a_{i1}, a_{i2}, \dots, a_{in})$ and all the other rows are the same as in Δ , and Δ_2 is a*

determinant whose i th row is equal to $(c_1, c_2, \dots, c_n)$ and all the other rows are the same as in Δ .

To prove this property, let us expand each of the three determinants Δ, Δ_1 and Δ_2 according to the i th row and note that in all the three determinants all the minors $\overline{M}_j^i$ of the elements of the i th row are the same. Hence, the formula $\Delta = \lambda\Delta_1 + \mu\Delta_2$ immediately follows from the equations $a_{ij} = \lambda b_j + \mu c_j$ ($j = 1, 2, \dots, n$).

The linear property is sure to be valid also for the case when the i th row is a linear combination not of two but of several rows. The linear properties is also valid for the columns of the determinant.

The three properties we have proved are the **principal** properties of a determinant revealing its nature.

The following five properties are logical consequences of the three principal properties.

Corollary 1. *A determinant with two similar rows (or columns) is equal to zero.* (Indeed, when two similar rows are interchanged, the determinant Δ does not change, on the one hand, and on the other hand, it changes sign to the opposite by virtue of property 2°. Thus $\Delta = -\Delta$, i.e. $2\Delta = 0$ or $\Delta = 0$.)

Corollary 2. *Multiplication of all elements of a certain row (or a certain column) of a determinant by the number λ is equivalent to the multiplication of the determinant by that number λ .*

In other words, *the common multiplier of all the elements of a certain row (or a certain column) of a determinant can be put before the sign of that determinant.* (This property follows from property 3° for $\mu = 0$.)

Corollary 3. *If all the elements of some row (or some column) of a determinant are zero, the determinant itself is equal to zero.* (This property follows from property 2° for $\lambda = 0$.)

Corollary 4. *If the elements of two rows (or two columns) of a determinant are proportional, the determinant is equal to zero.* (Indeed, by virtue of Corollary 2 the proportionality factor can be put before the sign of the determinant, and then a

determinant with two identical rows remains which is equal to zero in accordance with Corollary 1.)

Corollary 5. *If we add to the elements of some row (or some column) of a determinant the respective elements of another row (another column) multiplied by an arbitrary factor λ , the value of the determinant does not change.* (In fact, the determinant resulting from the indicated addition can be represented, by virtue of property 3°, as the sum of two determinants the first of which coincides with the original determinant and the second is equal to zero because of the proportionality of proportionality of the two rows (or columns) and Corollary 4.)

Remark. Corollary 5, as well as the linear property, admits of a more general formulation which we give for the rows: *if we add to the elements of row of a determinant the respective elements of the row which is a linear combination of several other rows of that determinant (with arbitrary coefficients) the value of the determinant does not change.*

Corollary 5 is widely used in actual calculations of determinants (the corresponding examples are given in 1.2.5).

Before formulating one more property of determinants, we introduce a useful notion of an algebraic adjunct (or a cofactor or a signed minor) of the given element of a determinant.

An algebraic adjunct of the given element a_{ij} of the n th-order determinant (1.11)

is a number equal to $(-1)^{i+j} \overline{M}_j^i$ and designated as A_{ij} .

Thus, the algebraic adjunct of a given element may differ from the minor of that element only in sign.

Using the notion of an algebraic adjunct, we can reformulate theorems 1.1 and 1.2 as follows: *the sum of the products of the elements of any row (any columns) of a determinant by the corresponding algebraic adjuncts of that row (column) is equal to that determinant.*

The corresponding formulas for the expansion of a determinant according to the i th row and the j th column can be rewritten as follows:

$$\Delta = \sum_{j=1}^n a_{ij}A_{ij} \text{ (for any } i=1,2,\dots,n), \quad (1.13')$$

$$\Delta = \sum_{i=1}^n a_{ij}A_{ij} \text{ (for any } j=1,2,\dots,n). \quad (1.21')$$

We can now formulate the property of a determinant.

4°. The property of algebraic adjuncts of adjacent rows (or columns). *The sum of the products of the elements of some row (or some column) of a determinant by the corresponding algebraic adjuncts of the elements of any other row (any other column) of a determinant by corresponding algebraic adjuncts of the elements of any other row (any other column) is equal to zero.*

We shall carry out the proof for the rows (it is quite analogous for the columns). Writing out formula (1.13')

$$\begin{vmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \cdot & \cdot & \cdot & \cdot \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{vmatrix} = A_{i1}a_{il} + A_{i2}a_{i2} + \dots + A_{in}a_{in} \quad (1.36)$$

we note that since the algebraic adjuncts $A_{i1}, A_{i2}, \dots, A_{in}$ are independent of the elements of the i th row $a_{i1}, a_{i2}, \dots, a_{in}$, equation (1.36) is an identity with respect to $a_{i1}, a_{i2}, \dots, a_{in}$ and is retained when the numbers $a_{i1}, a_{i2}, \dots, a_{in}$ are replaced by any other n numbers. Replacing $a_{i1}, a_{i2}, \dots, a_{in}$ by the respective elements of any (different from the i th) k th row $a_{k1}, a_{k2}, \dots, a_{kn}$, we get, on the left of (1.36), a determinant with two similar rows which is equal to zero in accordance with property 1. Thus,

$$A_{i1}a_{k1} + A_{i2}a_{k2} + \dots + A_{in}a_{kn} = 0$$

(for any noncoinciding i and k).

1.2.5. Examples of calculations of determinants.

When calculating determinants, we make wide use of formulas for expansion according to a row and a column and Corollary 5 which enables us, without

changing the value of the determinant, to add to any row (or columns) an arbitrary linear combination of its other rows (or columns). It is especially convenient to use the formula for expansion according to the rows (or columns) many of whose elements are zero. In particular, if only one element of the given row is different from zero, then the expansion according to that row contains only one term and reduces the problem of calculating a determinant of order n to that of calculating a determinant of order $n - 1$ (of the minor contained in the indicated term).

If, in the given row, several elements corresponding to the intersection of that row with several columns are different from zero, then, applying Corollary 5 to the indicated columns, we can turn into zero all the elements of that row, except one without changing the determinant.

Let us consider some examples.

Example 1. Suppose we have to calculate the following determinant of order 4:

$$\Delta = \begin{vmatrix} 4 & 99 & 83 & 1 \\ 0 & 8 & 16 & 0 \\ 60 & 17 & 143 & 20 \\ 15 & 43 & 106 & 5 \end{vmatrix}$$

Subtracting the trebled last column from the first column, we get

$$\Delta = \begin{vmatrix} 1 & 99 & 83 & 1 \\ 0 & 8 & 16 & 0 \\ 0 & 17 & 134 & 20 \\ 0 & 43 & 106 & 5 \end{vmatrix}$$

It is natural then to expand the determinant according to the first column. As a result, we get

$$\Delta = 1 \cdot \begin{vmatrix} 8 & 0 & 0 \\ 17 & 134 & 20 \\ 43 & 106 & 5 \end{vmatrix}.$$

Now, in the third-order determinant, we subtract the doubled first column from the second column and get

$$\Delta = \begin{vmatrix} 8 & 0 & 0 \\ 17 & 100 & 20 \\ 43 & 20 & 5 \end{vmatrix}.$$

Expanding now the last determinant of order 3 according to the first row, we finally get

$$\Delta = 8 \cdot \begin{vmatrix} 100 & 20 \\ 20 & 50 \end{vmatrix} = 8(500 - 400) = 800.$$

Example 2. Let us calculate the so-called *triangular determinant* whose all elements lying above the principal diagonal are zero:

$$\Delta_n = \begin{vmatrix} a_{11} & 0 & \dots & 0 & 0 \\ a_{21} & a_{22} & \dots & 0 & 0 \\ \dots & \dots & \dots & \dots & \dots \\ a_{(n-1)1} & a_{(n-1)2} & \dots & a_{(n-1)(n-1)} & 0 \\ a_{n1} & a_{n2} & \dots & a_{n(n-1)} & a_{nn} \end{vmatrix}.$$

Expanding the determinant Δ_n according to the last column, we find that it is equal to the product of the element a_{nn} by the $(n-1)$ th-order triangular determinant Δ_{n-1} equal to

$$\Delta_{n-1} = \begin{vmatrix} a_{11} & 0 & \dots & 0 \\ a_{21} & a_{22} & \dots & 0 \\ \dots & \dots & \dots & \dots \\ a_{(n-1)1} & a_{(n-1)2} & \dots & a_{(n-1)(n-1)} \end{vmatrix}.$$

We again expand the last determinant according to the last column and make sure that it is equal to the product of the element $a_{(n-1)(n-1)}$ by the $(n-2)$ th-order

triangular determinant Δ_{n-2} . Continuing similar arguments we arrive at the following expression for original determinant: $\Delta_n = a_{11}a_{22} \dots a_{nn}$.

Thus, *a triangular determinant is equal to the product of the elements lying on its principal diagonal.*

Remark 1. If all the elements of the determinant Δ lying below its principal diagonal are equal to zero, then the determinant is also equal to zero, then the determinant is also equal to the product of the elements lying on its principal diagonal (we can ascertain this using the scheme presented above, applying it, however, not the last columns but to the last rows, or we can simply transpose Δ and reduce this case to that considered above).

We can find, in a similar way, that a determinant whose all elements lying above (or below) the secondary diagonal are zero is equal to the product of the number $(-1)^{n(n-1)/2}$ by all the elements lying on that diagonal.

Example 3. The determinant of order $2n$ of the hypermatrix $\begin{vmatrix} A & O \\ B & C \end{vmatrix}$ in which A , B and C are arbitrary square matrices of order n and O is a null square matrix of order n , can serve as a generalization of a second-order triangular determinant. Let us ascertain that the following formula is valid for the indicated determinant:

$$\begin{vmatrix} A & O \\ B & C \end{vmatrix} = |A| |C| *. \quad (1.37)$$

(* Recall that we have agreed to use the designations $|A|$, $|B|$, $|C|$ for the determinants of the matrices A , B , C . . . respectively)

Applying Laplace's expansion, let us expand the determinant appearing on the left-hand side of (1.37) *according to the first n rows*. Since a determinant in which at least one column consists of zero is equal to zero it follows, according to (1.31), that only one term is different from zero, and the term (by virtue of the fact that $(-1)^{(1+\dots+n)+(1+\dots+n)} = 1$) is precisely equal to $|A| |C|$.

Remark 2. Using similar arguments, we can verify the validity of the formula

$$\begin{vmatrix} A & B \\ C & O \end{vmatrix} = (-1)^n |B| |C| \quad (1.38)$$

(A , B , C and O have the same sense as before).

For that purpose we must expand the determinant appearing on the left-hand side of (1.38) *according to the last n rows* and take into account that

$$(-1)^{[(n+1)+\dots+2n]+[1+\dots+n]} = (-1)^{2n(2n+1)/2} = (-1)^n.$$

Example 4. Let us now calculate the so-called **vandermonde determinant**

$$\Delta(x_1, x_2, \dots, x_n) = \begin{vmatrix} 1 & 1 & \dots & 1 \\ x_1 & x_2 & \dots & x_n \\ x_1^2 & x_2^2 & \dots & x_n^2 \\ \dots & \dots & \dots & \dots \\ x_1^{n-1} & x_2^{n-1} & \dots & x_n^{n-1} \end{vmatrix} \quad (1.39)$$

Subtracting the first column from all the rest columns, we get

$$\Delta(x_1, x_2, \dots, x_n) = \begin{vmatrix} 1 & 0 & \dots & 0 \\ x_1 & (x_2 - x_1) & \dots & (x_n - x_1) \\ x_1^2 & (x_2^2 - x_1^2) & \dots & (x_n^2 - x_1^2) \\ \dots & \dots & \dots & \dots \\ x_1^{n-1} & (x_2^{n-1} - x_1^{n-1}) & \dots & (x_n^{n-1} - x_1^{n-1}) \end{vmatrix}$$

It is natural now to expand the determinant according to the first row, as a result we get

$$\Delta(x_1, x_2, \dots, x_n) = \begin{vmatrix} (x_2 - x_1) & (x_3 - x_1) & \dots & (x_n - x_1) \\ (x_2^2 - x_1^2) & (x_3^2 - x_1^2) & \dots & (x_n^2 - x_1^2) \\ \dots & \dots & \dots & \dots \\ (x_2^{n-1} - x_1^{n-1}) & (x_3^{n-1} - x_1^{n-1}) & \dots & (x_n^{n-1} - x_1^{n-1}) \end{vmatrix}$$

Subtracting now from each row the preceding row multiplied by x_1 , we obtain

$$\Delta(x_1, x_2, \dots, x_n) = \begin{vmatrix} (x_2 - x_1) & (x_3 - x_1) & \dots & (x_n - x_1) \\ x_2(x_2 - x_1) & x_3(x_3 - x_1) & \dots & x_n(x_n - x_1) \\ . & . & . & . \\ x_2^{n-1}(x_2 - x_1) & x_3^{n-1}(x_3 - x_1) & \dots & x_n^{n-1}(x_n - x_1) \end{vmatrix}$$

Next we can put before the sign of the determinant the common factor of the first column equal to $(x_2 - x_1)$, the common factor of the second column equal to $(x_3 - x_1), \dots$, the common factor of the $(n-1)$ th. As a result we obtain

$$\Delta(x_1, x_2, \dots, x_n) = (x_2 - x_1)(x_3 - x_1) \dots (x_n - x_1) \Delta(x_2, x_3, \dots, x_n).$$

We do the same to the determinant $\Delta(x_2, x_3, \dots, x_n)$ on the right-hand side as to $\Delta(x_1, x_2, \dots, x_n)$. As a result we find that $\Delta(x_2, x_3, \dots, x_n) = (x_3 - x_2) \dots (x_n - x_2) \Delta(x_3, \dots, x_n)$.

Continuing similar arguments, we finally find that the original determinant (1.39) is equal to

$$\begin{aligned} \Delta(x_1, x_2, \dots, x_n) \\ = (x_2 - x_1)(x_3 - x_1) \dots (x_n - x_1)(x_3 - x_2) \dots (x_n - x_2) \dots \\ \dots (x_n - x_{n-1}). \end{aligned}$$

1.2.6. The determinant of the sum and the product of matrices.

It follows directly from the linear property of a determinant that *the determinant of the sum of two square matrices of the same order n , $A = \|a_{ij}\|$ and $B = \|b_{ij}\|$, is equal to the sum of all various determinants of order n which may result if we take a part of the rows (or columns) coinciding with the respective rows (or columns) of the matrix A and the other part coinciding with the respective rows (or columns) of B .*

Let us prove now that *the determinant of the matrix C , which is equal to the square matrix A by the square matrix B , is equal to the product of the determinants of the matrices A and B .*

Let the order of all the three matrices A , B and C be n , and let O be a zero square matrix of order n , and $(-1)I$ be the following matrix:

$$(-1)I = \begin{vmatrix} -1 & 1 & \dots & 0 \\ 0 & -1 & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & -1 \end{vmatrix}$$

By virtue of Example 2 of 1.2.5, the determinant of the matrix $(-1)I$ is equal to the number $(-1)^n$.

Let us consider the following two square hypermatrices of order $2n$:

$$\begin{vmatrix} A & O \\ (-1)I & B \end{vmatrix} \text{ and } \begin{vmatrix} A & C \\ (-1)I & O \end{vmatrix}.$$

By virtue of (1.37) and (1.38) of 1.2.5, the determinants of these matrices are

$$\begin{vmatrix} A & O \\ (-1)I & B \end{vmatrix} = |A| \cdot |B|, \quad \begin{vmatrix} A & C \\ (-1)I & O \end{vmatrix} = (-1)^n |(-1)I| \cdot |C| = (C).$$

Thus, it is sufficient to prove the equality of the determinants

$$\begin{vmatrix} A & O \\ (-1)I & B \end{vmatrix} \text{ and } \begin{vmatrix} A & O \\ (-1)I & B \end{vmatrix}.$$

In more detail, these two determinants can be written as follows:

$$\begin{vmatrix} a_{11} & a_{12} & \dots & a_{1n} & 0 & 0 & \dots & 0 \\ a_{21} & a_{22} & \dots & a_{2n} & 0 & 0 & \dots & 0 \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & a_{nn} & 0 & 0 & \dots & 0 \\ -1 & 0 & \dots & 0 & b_{11} & b_{12} & \dots & b_{2n} \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & a_{1n} & -1 & b_{21} & b_{22} & \dots & b_{nn} \end{vmatrix}, \quad (1.40)$$

$$\begin{vmatrix} a_{11} & a_{12} & \dots & a_{1n} & c_{11} & c_{12} & \dots & c_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} & c_{21} & c_{22} & \dots & c_{2n} \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & a_{nn} & c_{n1} & c_{n2} & \dots & c_{nn} \\ -1 & 0 & \dots & 0 & 0 & 0 & \dots & 0 \\ 0 & -1 & \dots & 0 & 0 & 0 & \dots & 0 \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & -1 & 0 & 0 & \dots & 0 \end{vmatrix}$$

To ascertain the equality of these two determinants, it is sufficient to note that the first n columns of these determinants coincide and every column of the second determinant (1.40) labelled $n+k$ (where $k=1, 2, \dots, n$) is obtained, by virtue of the

formula $c_{ij} = \sum_{k=1}^n a_{ik}b_{kj}$, by adding to the $(n+k)$ th column of the first determinant

(1.40) a linear combination of its first n columns with coefficients equal respectively to $b_{k1}, b_{k2}, \dots, b_{kn}$. Thus the determinants (1.40) are equal by virtue of the Corollary 5 of 1.2.3.

It should be pointed out in conclusion that it follows directly from formula (1.37)

that the determinant of the direct sum $|A \oplus B| = \begin{vmatrix} A & O \\ O & B \end{vmatrix}$ of two matrices A and B is

equal to the product of the determinants of those matrices.

1.2.7. The concept of an inverse matrix. Let A be a square matrix of order n and I , an identity square matrix of the same order (see 1.2).

The matrix B is called a **right inverse** matrix of the matrix A if $AB = I$.

The matrix C is called a **left inverse** matrix of the matrix A if $CA = I$.

Since both matrices A and I are square matrices of order n , the matrices B and C (provided that they exist) are also square matrices of order n .

Let us make sure that *if both matrices B and C exist, then they coincide*. Indeed, on the basis of equations (1.7) (see 1.2), the relations $AB = I, CA = I$ and the associative property of the product of matrices, we obtain

$$C = CI = C(AB) = (CA)B = IB = B.$$

Now a question arises about the condition concerning the matrix A under which both the left and the right inverse matrices of this matrix exist. (And obviously those matrices coincide.)

Theorem 1.4. *For the left and right inverse matrices of the matrix A to exist, it is necessary and sufficient that the determinant $\det A$ of the matrix A be nonzero.*

Proof. **(1) Necessity.** If at least one of the inverse matrices of the matrix A exists, say, B , then the relation $A \cdot B = I$ yields $\det A \cdot \det B = \det I = 1^{**}$, and hence $\det A \neq 0$. (** $\det I = 1$ by virtue of Example 2 in 1.2.5.)

(2) Sufficiency. Assume that the determinant $\Delta = \det A$ is nonzero. We shall use, as before, the designation A_{ij} for the algebraic adjuncts of the elements a_{ij} of the matrix A and form a matrix B whose i th row contains the algebraic adjuncts of the i th column of the matrix A divided by the value of the determinant Δ :

$$B = \begin{vmatrix} \frac{A_{11}}{\Delta} & \frac{A_{21}}{\Delta} & \dots & \frac{A_{n1}}{\Delta} \\ \frac{A_{12}}{\Delta} & \frac{A_{22}}{\Delta} & \dots & \frac{A_{n2}}{\Delta} \\ \cdot & \cdot & \cdot & \cdot \\ \frac{A_{1n}}{\Delta} & \frac{A_{2n}}{\Delta} & \dots & \frac{A_{nn}}{\Delta} \end{vmatrix}. \quad (1.41)$$

Let us make sure that this matrix B is both the right and the left inverse matrix of the matrix A .

It is sufficient to prove that the products AB and BA are identity matrices. To do that, it is sufficient to note that in both products *any off-diagonal element is equal to zero* since when the factor $1/\Delta$ is taken outside the matrix this element becomes equal to the sum of the products of the elements of one row (or one column) by the corresponding algebraic adjuncts of another row (or another column). As to the elements lying on the principal diagonal, all those elements in both products AB and BA are equal to unity because the sum of the products of the elements and the

corresponding algebraic adjuncts of one row (one column) is equal to the determinant. And thus, we have proved the theorem.

Remark 1. The square matrix A whose determinant $\det A$ is nonzero is customarily called a **non-singular** matrix.

Remark 2. In what follows we shall omit the terms “left” and “right” and speak of the *matrix B inverse of the non-singular matrix A* and defined by the relations $AB = BA = I$. It is also evident that the property of being an inverse matrix is reciprocal in the sense that if B is an inverse of A then A is an inverse of B . We shall use the designation A^{-1} for the inverse of the matrix A .

1.2 Theorem on the Base Minor of a Matrix

1.3.1. The concept of a linear dependence of rows. We have agreed to call the row* (* Each row can be regarded as a matrix, and, therefore, it is natural to use capital letters to designate rows.)

$A = (a_1, a_2, \dots, a_n)$ a linear combination of the rows $B = (b_1, b_2, \dots, b_n)$, $\dots$, $C = (c_1, c_2, \dots, c_n)$ if for some real numbers $\lambda, \dots, \mu$ the equations

$$a_j = \lambda b_j + \dots + \mu c_j \quad (j = 1, 2, \dots, n) \quad (1.42)$$

hold true. It is convenient to write these n equations as one equation

$$A = \lambda B + \dots + \mu C. \quad (1.43)$$

Every time we come across equation (1.43) we shall understand it in the sense of n equations (1.42).

We shall introduce now the notion of linear dependence of rows.

Definition. The rows $A = (a_1, a_2, \dots, a_n)$, $B = (b_1, b_2, \dots, b_n)$, $\dots$, $C = (c_1, c_2, \dots, c_n)$ are called **linearly dependent** if there are numbers $\alpha, \beta, \dots, \gamma$, not all equal to zero and such that the following equalities hold true:

$$\alpha a_j + \beta b_j + \dots + \gamma c_j = 0 \quad (j = 1, 2, \dots, n). \quad (1.44)$$

It is convenient to write n equations (1.44) as one equation

$$\alpha A + \beta B + \dots + \gamma C = O, \quad (1.45)$$

in which $O = (0, 0, \dots, 0)$ designates a zero row.

The rows which are not linearly dependent are called **linearly independent**. We can also give “its own” definition of linear independence of rows: *the rows $A, B, \dots, C$ are called **linearly independent** if equality (1.45) is only possible when all these numbers $\alpha, \beta, \dots, \gamma$ are equal to zero.*

We shall now prove a simple but very important statement.

Theorem 1.5. *For the rows $A, B, \dots, C$ to be linearly dependent it is necessary and sufficient that one of these rows be a linear combination of the other rows.*

Proof. (1) Necessity. Assume that the rows $A, B, \dots, C$ are linearly dependent, i.e. equality (1.45) in which at least one of the numbers $\alpha, \beta, \dots, \gamma$ is nonzero holds true. We assume, for definiteness, that $\alpha \neq 0$. Then, dividing (1.45) by α and introducing the designation $\lambda = -\beta/\alpha, \dots, \mu = -\gamma/\alpha$, we can rewrite (1.45) in the form

$$A = \lambda B + \dots + \mu C, \quad (1.46)$$

and this precisely means that the row A is a linear combination of the rows $B, \dots, C$.

(2) Sufficiency. Assume that one of the rows (say, A) is a linear combination of the other rows. Then, we can find numbers $\lambda, \dots, \mu$ such that equality (1.46) holds true. But this last equation can be rewritten as

$$(-1)A + \lambda B + \dots + \mu C = O^*.$$

Since one of the numbers $-1, \lambda, \dots, \mu$ is nonzero, the last equation establishes a linear dependence of the rows $A, B, \dots, C$. Thus we have proved the theorem.

As we have agreed, we can replace the term “rows” by the term “Columns” in all the arguments presented above.

1.3.2. Theorem on the base minor. Let us consider an arbitrary (not necessarily square) matrix

$$A = \begin{vmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_m & a_{m2} & \dots & a_{mn} \end{vmatrix} \quad (1.47)$$

* Here $O = (0, 0, \dots, 0)$ is a zero row.

The R th-order minor of the matrix A is a R th-order determinant with the elements lying in any R rows and any R columns of the matrix A (R is sure not to exceed the least of the numbers m and n).

We assume that at least one of the elements a_{ij} of the matrix A is nonzero. Then there is a positive integer r such that the following two conditions are satisfied: (1) the matrix A has a minor of order r different from zero, (2) every minor of the $(r+1)$ th or higher order (provided that they exist) is equal to zero.

We call the number r satisfying requirements (1) and (2) the rank of the matrix A . (* By definition, the rank of the matrix whose all elements are zeros is equal to zero.) An r th-order minor different from zero is known as a **base minor** (the matrix A can certainly have several r th-order minors different from zero). The rows and columns at whose intersection the base minor is located are called **base rows** and **base columns** respectively.

Let us prove the following fundamental theorem.

Theorem 1.6 (theorem on a base minor). *Base rows (base columns) are linearly independent. Any row (any column) of the matrix A is a linear combination of the base rows (base columns).*

Proof. We present all the arguments for rows.

If base rows were linearly dependent, then, in accordance with Theorem 1.5, one of these rows would be a linear combination of the other base rows and we could subtract the indicated linear combination of the other base rows and we could subtract the indicated linear combination from this row without changing the value of the base minor and get a row entirely consisting of zero and this would contradict the fact that the base minor is different from zero. Thus, the base rows are linearly independent.

Let us now prove that any row of the matrix A is a linear combination of the base rows. Since in arbitrary interchanges of rows (or columns) the determinant retains the property of being equal to zero, we can assume, without losing generality, that the base minor and the first r columns. Suppose j is any numbers from 1 to n and R is any numbers from 1 to m .

Let us ascertain that the determinant of order $r+1$

$$\begin{vmatrix} a_{11} & a_{12} & \dots & a_{1r} & a_{1j} \\ a_{21} & a_{22} & \dots & a_{2r} & a_{2j} \\ \cdot & \cdot & \cdot & \cdot & \cdot \\ a_{r1} & a_{r2} & \dots & a_{rr} & a_{rj} \\ a_{R1} & a_{R2} & \dots & a_{Rr} & a_{Rj} \end{vmatrix} \quad (1.48)$$

is equal to zero. If $j \leq r$ or $R \leq r$, then the indicated determinant is equal to zero because it has two identical columns or two identical rows.

Now if both numbers j and R exceed r , then (1.48) is a minor of the matrix A of order $(r+1)$ and any such minor is equal to zero (by the definition of a base minor). Thus, determinant (1.48) is equal to zero for all j from 1 to n and all R from 1 to m .

But then, expanding this determinant according to the last column and designating the cofactors of the elements of this column independent of the number j as $A_{1j} = c_1, A_{2j} = c_2, \dots, A_{rj} = c_r, A_{Rj} = c_{r+1}$, we find that

$$c_1 a_{1j} + c_2 a_{2j} + \dots + c_r a_{rj} + c_{r+1} a_{Rj} = 0$$

(for all $j = 1, 2, \dots, n$). Taking into account that in the last equations the cofactor $c_{r+1} = A_{Rj}$ coincides with the base minor which is obviously nonzero, we can divide each of these equations by c_{r+1} . But then, introducing the designations

$$\lambda_1 = -\frac{c_1}{c_{r+1}}, \lambda_2 = -\frac{c_2}{c_{r+1}}, \dots, \lambda_r = \frac{c_r}{c_{r+1}}$$

we find that $a_{Rj} = \lambda_1 a_{1j} + \lambda_2 a_{2j} + \dots + \lambda_r a_{rj}$ (for all $j = 1, 2, \dots, n$), and this precisely means that the R th row is a linear combination of the first r (base) rows. And this is what we set out to prove.

1.3.3. The necessary and sufficient condition for a determinant to be zero.

Theorem 1.7. *If the n th-order determinant Δ to be equal to zero, it is necessary and sufficient that its rows (columns) be linearly dependents.*

Proof. (1) Necessity. If the n th-order determinant Δ is equal to zero, then the base minor of its matrix has an order r obviously **smaller** than n . But then at least one row is non-base. By Theorem 1.6 that row is a linear combination of the base rows. We can include all the other rows in that linear combination putting zero before them.

Thus, one row is a linear combination of the other rows. But then, by Theorem 1.5, the rows of the determinant are linearly dependent.

(2) Sufficiency. If the rows of Δ are linearly dependent, then, by Theorem 1.5, one row A_i is a linear combination from the row A_j , we obtain a row consisting entirely of zero without changing the value of Δ . But then the determinant is equal to zero (by virtue of Corollary 3 from 2.4). We have proved the theorem.

Chapter – 2

LINEAR SPACES

The reader is sure to be acquainted from a course of analytic geometry with the operation of addition of free vectors and with the operation of multiplication of a vector by a real number, and also with the properties of these operations. In this chapter, we shall study sets of elements and that of multiplication of an element by a real number are defined in some way (no matter in what way), the indicated operations possessing the same properties as the respective operations of geometric vectors. Sets of this kind are known as linear spaces, they possess a number of common properties which will be established in the present chapter.

2.1 The Concept of a Linear Space

2.1.1. Definition of a linear space. *The set R of elements $x, y, z, \dots$ of any nature is called a linear (or affine) space if the following three requirements are met:*

*I. There is a rule according to which any two elements x and y of the set R are associated with a third element z of that set called the **sum** of the elements x and y and symbolized as $z = x + y$.*

*II. There is a rule according to which any element x of the set R and any real number λ are associated with an element u of that set called the **product of the element x by the number λ** and symbolized as $u = \lambda x$ or $u = x\lambda$.*

III. These two rules are subordinated to the following eight axioms.

1°. $x + y = y + x$ (commutative property of a sum).

2°. $(x + y) + z = x + (y + z)$ (associative property of a sum).

3°. *There is a zero element $\mathbf{0}$ such that $x + \mathbf{0} = x$ for any element x (a special role of a zero element).*

4°. *For every element x there is an **additive inverse** element x' such that $x + x' = \mathbf{0}$.*

5°. $\mathbf{1} \cdot x = x$ for every element x (a special role of the numerical factor 1).

6°. $\lambda(\mu x) = (\lambda\mu) x$ (associative property with respect to the numerical factor).

7°. $(\lambda + \mu)x = \lambda x + \mu x$ (distributive property with respect to the sum of the numerical factors).

8°. $\lambda(x + y) = \lambda x + \lambda y$ (distributive property with respect to the sum of the elements).

It should be emphasized that when we introduce the notion of a linear space, we abstract ourselves not only from the nature of the objects in question but also from the specific form of the rules of formation of a sum of elements and a product of an element by a number (it is only significant that those rules satisfy the eight axioms formulated in the definition given above).

Now if the nature of the objects in question and the form of the rules of formation of a sum of elements and a product of an element by numbers are indicated, then we call such a space a **specific** linear space.

Here are some examples of specific linear spaces.

Example 1. Let us consider a set of all free factors in a three-dimensional space. We define the operations of addition of the indicated vectors and their multiplication by scalars as is usually done in courses of analytic geometry (we define the addition of vectors by the “parallelogram” rule; in the multiplication of a vector by a real number λ , the length of the vector is multiplied by $|\lambda|$ and the direction remains unchanged when $\lambda > 0$, and changed to the opposite when $\lambda < 0$).

It is a simple practice to verify the validity of axioms 1°-8°. Thus, the set of all free vectors in a space with the operations of addition and their multiplication by scalars so defined is a linear space which will be designated as B_3 .

Similar sets of vectors on a plane and on a straight line, which are also linear spaces, will be denoted by B_2 and B_1 respectively.

Example 2. Let us consider a set $\{x\}$ of all **positive** real numbers. We define the sum of two elements x and y of that set as a product of real numbers x and y (understood in the same sense as in the theory of real numbers). The product of the

element x of the set $\{x\}$ is a real number 1 and the additive inverse element (for the given element x) is the real number $1/x$.

The validity of all axioms 1°-8° is easy to verify. Indeed, the validity of axioms 1° and 2° follows from the commutative and associative properties of a product of real numbers; the validity of axioms 3° and 4° follows from the simple equations $x \cdot 1 = x$, $x \cdot \frac{1}{x} = 1$ (for any real $x > 0$), axiom 5° is equivalent to the equation $x^1 = x$, axioms 6° and 7° are valid because for any $x > 0$ and any real λ and μ the relations $(x^\mu)^\lambda = x^{\lambda\mu}$, $x^{\lambda+\mu} = x^\lambda \cdot x^\mu$ hold true, and, finally the validity of axioms 8° follows from the fact that the equality $(xy)^2 = x^\lambda y^\lambda$ holds true for any positive x and y and for any real λ . Thus we have ascertained that the set $\{x\}$ with the operations of addition of elements and their multiplication by numbers so defined is a linear space.

Example 3. The set A^n whose elements are ordered collections of n arbitrary real numbers $(x_1, x_2, \dots, x_n)$ serves as an important example of a linear space. We denote the elements of this set by one symbol $\mathbf{x}$, i.e. we write $\mathbf{x} = (x_1, x_2, \dots, x_n)$ and call the real numbers $x_1, x_2, \dots, x_n$ the **coordinates** of the element $\mathbf{x}$.

In mathematical analysis it is customary to call the set A^n an **n - dimensional coordinates space**. In algebraic interpretation, we can regard A^n as a collection of various rows each of which contains n real numbers (we did just that in 1.3).

The following rules define the operations of addition of the elements of the set A^n and their multiplication by real numbers:

$$(x_1, x_2, \dots, x_n) + (y_1, y_2, \dots, y_n) = (x_1 + y_1, x_2 + y_2, \dots, x_n + y_n),$$

$$\lambda(x_1, x_2, \dots, x_n) = (\lambda x_1, \lambda x_2, \dots, \lambda x_n).$$

We invite the reader to verify the validity of axioms 1°-8° and of the fact that the element $\mathbf{0} = (0, 0, \dots, 0)$ is an additive inverse element of $(x_1, x_2, \dots, x_n)$.

Example 4. Let us now consider the set $C[a, b]$ of all functions $x = x(t)$ defined and continuous on the interval $a \leq t \leq b$. The operations of addition of such

functions and their multiplication by real numbers are defined by ordinary rules of mathematical analysis. The validity of axioms 1° – 8°, which can be easily verified, enables us to infer that the set $C[a, b]$ is a linear space.

Example 5. We can consider as an example of a linear space the set $\{P_n(t)\}$ of all algebraic polynomials of a degree not exceeding a natural number n , with the operations defined as in the preceding example. Note that if the set $\{P_n(t)\}$ is considered on the interval $a \leq t \leq b$, then it is a **subset** of the linear space $C[a, b]$ considered in Example 4.

Remark 1. To explain the notion of a linear space in more detail, we give examples of sets which are not linear spaces for some reason or other:

- (a) the set of all vectors of a space except for vectors collinear with some straight-line l (since we cannot add together the vectors symmetric about the indicated line l within that set);
- (b) the set of all polynomials of a degree exactly equal to the number n (the sum of two polynomials of this kind may turn out to be a polynomial of a degree lower than n);
- (c) the set of all polynomials of a degree not exceeding a natural n , whose all coefficients are positive (the elements of such a set cannot be multiplied by negative real numbers).

Remark 2. Note that the elements of an arbitrary linear space are customarily called **vectors**. The fact that the term “vector” is often used in a narrower sense does not lead to a misunderstanding. On the contrary, appealing to familiar geometric notions, we can elucidate and often even foresee a number of results valid for linear spaces of arbitrary nature.

Remark 3. In the definition of a linear space we have formulated we took the numbers $\lambda, \mu, \dots$ from the set of **real** numbers. Therefore, it is natural to call the space we have defined a **real linear space**. In a wider approach, we can take $\lambda, \mu, \dots$ from the set of **complex** numbers. In that case, we arrive at the notion of a **complex linear space**.

2.1.2. Some properties of arbitrary linear space. As logical consequences of axioms 1° - 8°, we can obtain a number of statements valid for arbitrary linear spaces. We shall consider two statements as an example.

Theorem 2.1. *There is a single zero element in an arbitrary linear space and there is a single additive inverse element of every element x .*

Proof. The existence of **at least one** zero element is stated in axiom 3°. Let us assume that there are two zero elements) 0_1 and 0_2 . Then, setting first $x=0_1, 0=0_2+0_1=0_2$ in axiom 3° and then $x=0_2, 0=0_1$ we obtain two equations $0_1+0_2=0_1, 0_2+0_1=0_2$ whose left-hand sides are equal by virtue of axiom 1°. This means that the right-hand sides of the two last equations are equal because of the transitivity of the equality sign, i. e. $0_1=0_2$, and we have thus established the uniqueness of the zero element.

Axiom 4° asserts the existence of **at least one** additive inverse element y of the element x . Let us assume that there are two additive inverse elements y_1 and y_2 of a certain element x such that $x+y_1=0$ and $x+y_2=0$. But then, by virtue of axioms 3°, 2° and 1°, $y_1=y_1+0=y_1+(x+y_2)=(y_1+x)+y_2=0+y_2+0=y_2$, i.e. $y_1=y_2$, and we have proved the uniqueness of the additive inverse element of any element x . We have proved the theorem.

Theorem 2.2. *In an arbitrary linear space*

(1) *the zero element 0 is equal to the product of the arbitrary element x by a real number 0 ,*

(2) *the additive inverse element of every element x is equal to the product of the element x by a real number -1 .*

Proof. (1) Suppose x is an arbitrary element any y is its additive inverse. Applying successively axioms 3°, 4°, 2°, 5°, 1°, and 7° and again 5° and 4°, we find that

$$x \cdot 0 = x \cdot 0 + 0 = x \cdot 0 + (x + y) = (x \cdot 0 + x) + y = (x \cdot 0 + x \cdot 1) + y = x(0+1) + y = x \cdot 1 + y = x + y = 0, \text{ i.e. } x \cdot 0 = 0.$$

(2) Suppose x is an arbitrary element and $y = (-1) \cdot x$. Using axioms $5^0, 7^0$ and 1^0 and the equality $x \cdot 0 = 0$ we have proved, we get an equality $x + y = x + (-1)x = 1 \cdot x + (-1) \cdot x = [1 + (-1)] \cdot x = 0 \cdot x = x \cdot 0 = 0$ which proves (by virtue of axiom 4^0) that y is an additive inverse element of x . And that was the proof of the theorem.

We must note in conclusion that axioms $1^0 - 4^0$ make it possible to prove the existence and uniqueness of the difference between any two element x and y of a linear space which is defined as an element z satisfying the condition $z + y = x^*$. (*It is sufficient to repeat the proof given in the theory of real numbers in *Fundamentals of Mathematical Analysis*, Ilyin and Poznyak) (The sum $z = x + y(-1)y$ is precisely that element.)

2.2. The Basis and the Dimensionality of a Linear space

2.2.1. The concept of a linear dependence of the elements of a linear space. In the course of analytic geometry we introduced the notion of a linear dependence of vectors, and in 1.3.1 of the present book we introduced the notion of a linear dependence of rows (or, what is the same thing, of the elements of the space A^n considered in Example 3 in 2.1.1).

A generalization of these notions is the notion of linear dependence of elements of an arbitrary linear space which we shall now consider.

Let us consider an arbitrary real linear space R with elements $x, y, \dots, z, \dots$

A **linear combination** of the elements $x, y, \dots, z$ of the space R is the sum of the products of these elements by arbitrary real numbers, i.e. an expression of the form

$$\alpha x + \beta y + \dots + \gamma z, \quad (2.1)$$

where $\alpha, \beta, \dots, \gamma$ are any real numbers.

Definition 1. The elements $x, y, \dots, z$ of the space R are said to be **linearly dependent** if there are real numbers $\alpha, \beta, \dots, \gamma$ at least one of which is different from zero, such that the linear combination of the elements $x, y, \dots, z$ with the indicated numbers is a zero element of the space R , i.e. there holds an equality

$$\alpha x + \beta y + \dots + \gamma z = 0. \quad (2.2)$$

The elements $x, y, \dots, z$ which are not linearly dependent are said to be **linearly independent**.

Here is another definition of linearly independent vectors based on the logical negation of content of Definition 1.

Definition 2. *The elements $x, y, \dots, z$ of the space R are said to be **linearly independent** if the linear combination (2.1) is a zero element of the space R only under the condition $\alpha = \beta = \dots = \gamma = 0$.*

Theorem 2.3. *For the elements $x, y, \dots, z$ of the space R to be linearly dependent, it is necessary and sufficient that one of these elements be a linear combination of the other elements.*

Proof. (1) Necessity. Suppose the elements $x, y, \dots, z$ are linearly dependent, i.e. equality (2.2), in which at least one of the numbers $\alpha, \beta, \dots, \gamma$ is nonzero, holds true. Assume, for definiteness, $\alpha \neq 0$. Then, dividing (2.2) by α and introducing the designations $\lambda = -\frac{\beta}{\alpha}, \dots, \mu = -\frac{\gamma}{\alpha}$, we can rewrite (2.2) in the form

$$x = \lambda y + \dots + \mu z, \quad (2.3)$$

and this means that the elements x is a linear combination of the elements $y, \dots, z$.

(2) Sufficiency. Assume that one of the elements (say, x) is a linear combination of the other elements. Then, there can be found numbers $\lambda, \dots, \mu$ such that equality (2.3) holds true. But this last equality can be rewritten as

$$(-1)x + \lambda y + \dots + \mu z, \quad (2.4)$$

Since one of the numbers $(-1), \lambda, \dots, \mu$ is different from zero, equation (2.4) establishes a linear dependence of the elements $x, y, \dots, z$. And that is what we had to prove.

The following two simple statements hold true.

1. If there is a zero elements among the elements among the elements $\mathbf{x}, \mathbf{y}, \dots, \mathbf{z}$, then these elements are linearly dependent. Indeed, if say, $\mathbf{x} = \mathbf{0}$, then equality (2.2) is valid for $\alpha = 1, \beta = \dots = \gamma = 0$.

2. If a part of the elements $\mathbf{x}, \mathbf{y}, \dots, \mathbf{z}$ are linearly dependent, then all these elements are linearly dependent. Indeed, if, say, the elements $\mathbf{y}, \dots, \mathbf{z}$ are linearly dependent, then the equality $\beta\mathbf{y} + \dots + \gamma\mathbf{z} = \mathbf{0}$, in which not all the numbers $\beta, \dots, \gamma$ and with $\alpha = 0$ holds true.

We conclude with studying the linear dependence of the elements of the space A^n introduced in Example 3 in 2.1.1.

We shall prove that n elements of the indicated space

$$\left. \begin{aligned} \mathbf{e}_1 &= (1, 0, 0, \dots, 0), \\ \mathbf{e}_2 &= (0, 1, 0, \dots, 0), \\ \mathbf{e}_3 &= (0, 0, 1, \dots, 0), \\ &\vdots \\ \mathbf{e}_n &= (0, 0, 0, \dots, 1) \end{aligned} \right\} \quad (2.5)$$

are linearly independent and the collection of n elements (2.5) and one more **arbitrary** element $\mathbf{x} = (x_1, x_2, \dots, x_n)$ of the space A^n forms a linearly dependent system of elements.

Let us consider the linear combination of elements (2.5) with any numbers $\alpha_1, \alpha_2, \alpha_3, \dots, \alpha_n$. By virtue of the axioms, this linear combination is an element

$$\alpha_1, \alpha_2 \mathbf{e}_2 + \dots + \alpha_n \mathbf{e}_n = (\alpha_1, \alpha_2, \dots, \alpha_n),$$

which is a zero element only under the condition $\alpha_1 = \alpha_2 = \dots = \alpha_n = 0$. But this precisely means a linear independence of elements (2.5).

Let us now prove that a system consisting of n elements (2.5) and one more **arbitrary** element $\mathbf{x}(x_1, x_2, \dots, x_n)$ of the space A^n is linearly dependent. By virtue of Theorem 2.3, it is sufficient to prove that that *element* $\mathbf{x} = (x_1, x_2, \dots, x_n)$ is a

linear combination of elements (2.5) and this is obvious since by virtue of the axioms

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

2.2.2. The basis and the coordinates. Let us consider an arbitrary real linear space R .

Definition. A collection of linearly independent elements $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ of the space R is said to be the **basis** of that space if for every element $\mathbf{x}$ of the space R there are real numbers $x_1, x_2, \dots, x_n$ such that the following equality holds true:

$$\mathbf{x} = x_1 \mathbf{e}_1 + x_2 \mathbf{e}_2 + \dots + x_n \mathbf{e}_n. \quad (2.6)$$

In the case, equation (2.6) is called the **resolution of the element $\mathbf{x}$ into components** $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ and the numbers $x_1, x_2, \dots, x_n$ are called the **coordinates** of the element $\mathbf{x}$ (with respect to the basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$).

Let us prove that *every element $\mathbf{x}$ of the linear space R can be resolved into the components $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ in only one way*, i.e. the coordinates of every element $\mathbf{x}$ with respect to the basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ are defined **uniquely**.

Let us assume that besides resolution (2.6) one more resolution into the same components

$$\mathbf{x} = x'_1 \mathbf{e}_1 + x'_2 \mathbf{e}_2 + \dots + x'_n \mathbf{e}_n \quad (2.7)$$

is possible for a certain element $\mathbf{x}$.

A term-by-term subtraction of equations (2.6) and (2.7) leads to a relation*

$$(x_1 - x'_1) \mathbf{e}_1 + (x_2 - x'_2) \mathbf{e}_2 + \dots + (x_n - x'_n) \mathbf{e}_n = \mathbf{0}. \quad (2.8)$$

(*The possibility of a term-by-term subtraction of equations (2.6) and (2.7) and the grouping of the terms we perform follows axioms 1° - 8°.)

Because of the linear independence of the base elements $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$, relation (2.8) leads to equations $x_1 - x'_1 = 0, x_2 - x'_2 = 0, \dots, x_n - x'_n = 0$ or $x_1 = x'_1, x_2 = x'_2, \dots, x_n = x'_n$. the uniqueness of resolution into components. The

practical significance of the basis lies also in the fact that when it is specified the operations of addition of elements and their multiplication by numbers become the corresponding operations on numbers, the coordinates of those elements. The following assertion holds true.

Theorem 2.4 *When any two elements of the linear space R are added together, their coordinates (with respect of any basis of the space R) are added; when an arbitrary element is multiplied by any number λ , all the coordinates of that element are multiplied by λ .*

Proof. Suppose $e_1, e_2, \dots, e_n$ is a arbitrary basis of the space R , $x = x_1e_1 + x_2e_2 + \dots + x_ne_n$ and $y = y_1e_1 + y_2e_2 + \dots + y_ne_n$ are any two elements of that space.

Then, by virtue of axioms 1°-8°, we have

$$x + y = (x_1 + y_1)e_1 + (x_2 + y_2)e_2 + \dots + (x_n + y_n)e_n,$$

$$\lambda x = (\lambda x_1)e_1 + (\lambda x_2)e_2 + \dots + (\lambda x_n)e_n.$$

This completes the proof since the resolution into components is unique.

Here are some examples of the bases of specific linear spaces. It is known from analytic geometry that any three noncoplanar vectors form a basis in the linear space B_3 of all free vectors (this space was considered in example 1 in 2.1.1).

Now note that the collection of n elements (2.5) considered at the end of 2.2.1 forms a basis in the linear space A^n introduced in Example 3 in 2.1.1.

In fact, we have proved at the end of 2.2.1 that elements (2.5) are linearly independent and that any element $x = (x_1, x_2, \dots, x_n)$ of the space A^n is a linear combination of elements (2.5).

Let us finally ascertain that the basis of the linear space $\{x\}$ introduced in Example 2 in 2.1.1, consists of one element which can be any **nonzero** elements of that space (i.e. any positive real number x_0 not equal 1). It is sufficient to prove that for any positive real number x there is a real number λ such that $x = x_0^\lambda$. (* Recall that the

product of the element $\mathbf{x}_0$ by the number λ is defined as the number x_0^λ .) But this is obvious: it is sufficient to take $\lambda = \log_{x_0} x$.

2.2.3 The dimensionality of a linear space. As before, we shall consider an arbitrary real linear space R .

Definition 1. *The linear space R is called **n -dimensional** if there are n linearly independent elements in it whereas any $(n+1)$ elements are linearly dependent. The number n is called a **dimension** of the space R .*

The dimension of the space R is usually designated as $\dim R$.

Definition 2. *The linear space R is said to be **infinite-dimensional** if it contains any number of linearly independent elements.* **(** The fact that the space R is infinite-dimensional is symbolized as $\dim R = \infty$.)

In this book we shall study only spaces of finite dimension n . Infinite-dimensional spaces form a subject requiring special investigation. (They are studied in chapters 10 ad 11 of *Fundamentals of Mathematical Analysis*, Part 2.)

Let us elucidate the connection between the notion of dimensionality of a space and that of a basis introduced in 2.2.2.

Theorem. 2.5. *If R is a linear space of dimension n , then any n linearly independent elements of that space form its basis.*

Proof. Suppose $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ is any system of n linearly independent elements of the space R (the existence of a least one system of that kind follows from Definition 1).

If $\mathbf{x}$ is any elements of R , then, according to Definition 1, a system of $(n+1)$ elements $\mathbf{x}, \mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ is linearly dependent, i.e. there are numbers $\alpha_0, \alpha_1, \alpha_2, \dots, \alpha_n$ which are not all equal to zero and such that the equality

$$\alpha_0 \mathbf{x} + \alpha_1 \mathbf{e}_1 + \alpha_2 \mathbf{e}_2 + \dots + \alpha_n \mathbf{e}_n = 0 \quad (2.9)$$

holds true. Note that the number α_0 is obviously nonzero (since otherwise equation (2.9) would yield a linear dependence of the elements $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$). But then,

dividing equation (2.9) by α_0 setting $x_1 = -\frac{\alpha_1}{\alpha_0}, x_2 = -\frac{\alpha_2}{\alpha_0}, \dots, x_n = -\frac{\alpha_n}{\alpha_0}$, we get $x = x_1 e_1 + x_2 e_2 + \dots + x_n e_n$ (2.10) form (2.9). Since x is an arbitrary element of R , equation (2.10) proves that the system of elements $e_1, e_2, \dots, e_n$ is a basis of the space R . We have proved the theorem.

Theorem 2.6. *If the linear space R has a basis consisting of n elements, then the dimension of R is n .*

Proof. Suppose the system of n elements $e_1, e_2, \dots, e_n$ is a basis of the space R . It is sufficient to prove that **any** $(n+1)$ elements $x_1, x_2, \dots, x_{n+1}$ of that space are linearly dependent. * (* Because the basis elements $e_1, e_2, \dots, e_n$ form a system of n linearly independent elements of the space R .)

Resolving each of these elements into components, we get

$$x_1 = a_{11}e_1 + a_{12}e_2 + \dots + a_{1n}e_n,$$

$$x_2 = a_{21}e_1 + a_{22}e_2 + \dots + a_{2n}e_n,$$

$$\dots \dots \dots$$

$$x_{n+1} = a_{(n+1)1}e_1 + a_{(n+1)2}e_2 + \dots + a_{(n+1)n}e_n,$$

where $a_{11}, a_{12}, \dots, a_{(n+1)n}$ are some real numbers.

The linear dependence of the elements $x_1, x_2, \dots, x_{n+1}$ is evidently equivalent to the linear dependence of the rows of the matrix

$$A = \begin{vmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{(n+1)1} & a_{(n+1)2} & \dots & a_{(n+1)n} \end{vmatrix}.$$

But the rows of this matrix are obviously linearly dependent since the order of the base minor of this matrix (containing $n+1$ rows and n columns) does not exceed n and at least one of its $n+1$ rows is not a base row and is, by the theorem on a base

minor, a linear combination of base (and, evidently, all the other) rows. We have proved the theorem.

Directing our attention to the examples considered at the end of 2.2.2, we can say that the dimension of the space B_3 of all free vectors is 3, that of the space A^n is n and the dimension of the space $\{x\}$ is unity.

The linear space $C[a, b]$ of all functions $x = x(t)$, defined and continuous on the interval $a \leq t \leq b$, can serve as an example of an infinite-dimensional space (see Example 4 in 2.1.1).

Indeed, for any number n , a system of $(n+1)$ elements $1, t, t^2, \dots, t^n$ of that space is linearly independent (since otherwise a certain polynomial $C_0 + C_1t + C_2t^2 + \dots + C_nt^n$, in which not all the coefficient $C_0, C_1, \dots, C_n$ are equal to zero, would be identically equal to zero on the interval $a \leq t \leq b$).

2.2.4 The concept of isomorphism of linear spaces. We show here that various linear spaces of one and the same dimension n , in the sense of the properties connected with the operations introduced in those spaces, do not differ essentially from one another.

Since we have introduced in linear spaces only the operations of addition of elements and their multiplication by numbers, it is natural to formulate the following definition.

Definition. *Two arbitrary real linear space R and R' are said to be **isomorphic** (or; more precisely, linearly **isomorphic**) to each other if a one-to-one correspondence* can be established between the elements of those spaces so that if the elements x and y of the space R are associated with the respective elements x' and y' of the space R' , then the elements $x + y$ is associated with the element $x' + y'$ and the element λx is associated with the element $\lambda x'$ for any real λ .*

* Recall that the correspondence between the elements of two sets R and R' is said to be **one-to-one** if that correspondence relates every element of R to one and only one element of R' . In that case every element of R' corresponds to one and only one element of R .

Note that *if the linear space R is associated with a zero element of R' and vice versa.* (Indeed, suppose the element x of R is associated with a certain element x' and R' . Then the element $0.x$ of R is associated with the element $0.x'$ of R' .)

It follows that if in the case of isomorphism the elements $x, y, \dots, z$ of the space R correspond to the respective elements $x', y', \dots, z'$ of the space R' , then the linear combination $\alpha x + \beta y + \dots + \gamma z$ is zero element of the space R if and only if the space R' .

But this means that *if the space R and R' are isomorphic, then the maximum number of linearly independent elements in each of these spaces is the same.*

In other words, *two isomorphic spaces must have the same dimension.*

Hence, *spaces of different dimensions cannot be isomorphic.*

Let us now prove the following statement.

Theorem 2.7. *Any two n -dimensional real linear spaces R and R' are isomorphic.*

Proof. Let us choose some basis $e_1, e_2, \dots, e_n$ in R and some basis $e'_1, e'_2, \dots, e'_n$ in R' . We put every element $x' = x_1 e'_1 + x_2 e'_2 + \dots + x_n e'_n$ of the space R (i.e. we take as x' the element of R' which has *the same coordinates* with respect to the basis $e_1, e_2, \dots, e_n$).

We shall verify that the correspondence we have established is one-to-one. Indeed, every element x of the spaces R is in a one-to-one correspondence with the coordinates $x_1, x_2, \dots, x_n$, which, in their turn, define a single element x' of the space R' . Because of the equivalence of the space R and R' , every element x' of the space R' is associated with only one element x of the space R .

It remains to point out that if the elements x and y of R are put into correspondence with the respective elements x' and y' of R' , then, by virtue of Theorem 2.4, the element $x + y$ is associated with the elements $x' + y'$ and the element λx with the element $\lambda x'$. We have proved the theorem.

It follows from the consideration presented above that the only essential characteristic of a finite-dimensional linear space is its dimensionality.

2.3 Subspace of Linear Spaces

2.3.1. The concepts of a subspace and a linear closure. Let us assume that a subset L of the linear space R meets the following **two requirements**.

1°. *If the elements x and y belong to the subset L , then the sum $x + y$ also belongs to the subset.*

2°. *If the element x belongs to the subset L and λ is any real number, then the element λx also belongs to the subset L .*

Let us make sure that the subset L meeting requirements 1° and 2° is itself a linear space. It is sufficient to verify the validity of axioms 1°-8° from the definition of a linear space for the elements of the subset L . All the indicated axioms, except for 3° and 4° are obviously valid for the elements of the subset L since they are valid for **all** elements of the space R . It remains to verify the validity of axioms 3° and 4°. Suppose x is any element of the subset L and λ is any real number. Then, by virtue of requirement 2°, the element λx also belongs to L . It remains to note that (according to Theorem 2.2) for $\lambda = 0$ this element λx turns into a zero element of the space R and for $\lambda = -1$ it turns into an additive inverse element of x . Thus, the zero element and the additive inverse element (of every x) belong to the subset L and this means that axioms 3° and 4° are valid for the elements of the subset L . We have thus completely proved that the subset L is itself a linear space.

Definition. *The subset L of the linear space R meeting requirements 1° and 2° is called a **linear subspace** (or simply **subspace**) of the space R .*

Here are some examples of subspaces:

(1) the so-called **null subspaces**, i.e. a subset of the linear space R consisting of one zero element,

(2) the whole space R (which can certainly be regarded as a subspace).

Both these subspaces are customarily called **improper**.

Here are some more interesting examples of subspaces.

Example 1. The subset $\{P_n(t)\}$ of all algebraic polynomials of a degree not exceeding a natural number n , in the linear space $C[a, b]$ of all functions $x = x(t)$

defined and continuous on the interval $a \leq t \leq b$ ^{**} (the validity of requirements 1° and 2° for the elements of the subset $\{P_n(t)\}$ is obvious). (** The space $C[a, b]$ was introduced in Example 4 in 2.1.1.)

Example 2. The subset B_3 of all free vectors, parallel to some plane, in the linear space B_3 of all free vectors (the validity of requirements 1° and 2° for the elements of B_2 is obvious).

Example 3. Assume that $x, y, \dots, z$ is a collection of elements of a certain linear space R . The **linear closure** of the elements $x, y, \dots, z$ is a collection of all linear combinations of those elements, i.e. the set of elements of the form

$$\alpha x + \beta y + \dots + \gamma z,$$

where $\alpha, \beta, \dots, \gamma$ are any real numbers.

We shall agree to use the symbol $L(x, y, \dots, z)$ for the linear closure of the elements $x, y, \dots, z$.

Requirements 1° and 2° formulated at the beginning of this section are evidently satisfied for the linear closure of the arbitrary elements $x, y, \dots, z$ of the linear space R . Therefore, *every linear closure is a subspace of the principal linear space R* . This subspace evidently contains the elements $x, y, \dots, z$ must also contain all linear combinations of those elements. Therefore, *the linear closure of the elements $x, y, \dots, z$ is the least subspace containing the elements $x, y, \dots, z$* .

The linear closure of the elements $1, t, t^2, \dots, t^n$ of the linear space $C[a, b]$ of all function $x = x(t)$ defined and continuous on the interval $a \leq t \leq b$ can serve as an example of a linear closure. This linear closure is evidently a set $\{P_n(t)\}$ of all algebraic polynomials of a degree not exceeding n .

Some other examples of subspaces will be studied in 2.3.3.

Let us now consider the dimensionality of a subspace (of a linear closure, in particular).

We can state that *the dimension of any subspace of an n -dimensional linear space R does not exceed the dimension n of the space R* (since every linearly independent system of elements of a subspace is at the same time a linearly independent system of elements of the entire space R).

To be more precise, we can state that *if the subspace L does not coincide with the whole n -dimensional linear space R , then the dimension of L is strictly less than n .*

This follows from the fact that if the dimensions of L and R are both equal to n , then every basis of the subspace L is (by virtue of Theorem 2.5) the basis of the whole space R since it consists of n elements.

Not that if a basis $e_1, e_2, \dots, e_n$ has been chosen in the whole space R , then the base elements of the subspace L cannot, in general, be chosen from the elements $e_1, e_2, \dots, e_n$ (since, in a general case, none of the elements $e_1, e_2, \dots, e_n$ necessarily belong to L). The converse **statement** is true, however: *if the elements $e_1, e_2, \dots, e_n$ form the basis of a k -dimensional subspace of an n -dimensional linear space R , then this basis can be complemented by the elements $e_{k+1}, \dots, e_n$ of the space R so that the collection of the elements $e_1, \dots, e_k, e_{k+1}, \dots, e_n$ forms a basis of the whole space R .*

Let us prove this statement. If $k < n$, then there is an element e_{k+1} of the space R such that the elements $e_1, e_2, \dots, e_k, e_{k+1}$ are linearly independent (otherwise, the space R would be k -dimensional). Furthermore, if $k + 1 < n$, then there is an element e_{k+2} of the space R such that the elements $e_1, e_2, \dots, e_k, e_{k+1}, e_{k+2}$ are linearly independent (otherwise, the Space R would be $(k + 1)$ -dimensional). Continuing a similar argument, we shall prove the statement we have formulated.

In conclusion, we shall prove a significant theorem on the dimensionality of a linear closure.

Theorem 2.8 *The dimension of the linear closure $L(x, y, \dots, z)$ of the elements $x, y, \dots, z$ is equal to maximum number of linearly independent elements in the system of elements $x, y, \dots, z$. In particular, if the elements $x, y, \dots, z$ are linearly independent, then the dimension of the linear closure $L(x, y, \dots, z)$.*

Proof. We assume that there are r linearly independent elements among the elements $x, y, \dots, z$ (we designate them as $x_1, x_2, \dots, x_r$), and any $(r+1)$ elements from the elements $x, y, \dots, z$ is a certain linear combination of the elements $x_1, x_2, \dots, x_r$, and since, by definition, every element of the linear closure $L(x, y, \dots, z)$ is a certain linear combination of elements $x, y, \dots, z$ every element of the indicated linear closure is a certain linear combination of the elements $x_1, x_2, \dots, x_r$ alone. But this precisely means that the system of linearly independent elements $x_1, x_2, \dots, x_r$ form a basis of the linear closure $x, y, \dots, z$ and the dimensionality of $x, y, \dots, z$ is equal to r . Thus, the proof is complete.

* This fact can be established with the aid of the same arguments which were presented in the proof of Theorem 2.5.

2.3.2 A new definition of a rank of a matrix. In 1.3 we defined the rank of the arbitrary matrix A as the **order of its base minor**, i.e. as the number r satisfying the requirement that the matrix A should have a nonzero minor of order r and should not have nonzero minors of order exceeding r .

We shall prove in this subsection that *the rank of the arbitrary matrix A is equal to the maximum number of linearly independent rows (or column) of that matrix*.

This statement yields a new definition of the rank of a matrix as a maximum number of linearly independent rows (or columns) of that matrix*.

We present all the arguments for the rows (since they are similar for the columns). We consider the linear closure of the base rows of an arbitrary matrix A , containing m' rows and n columns, in the linear space A^n (introduced in Example 3 in 2.1.1) and assume that the number of base rows is equal to r . It follows from Theorem 1.6 on a base minor that any row of the linear independence of r base rows and from Theorem 2.8 it follows that the dimension of the indicated linear closure is equal to r . Hence any $(r+1)$ elements of that linear closure (and, in particular, any $(r+1)$ rows of the matrix A) are linearly dependent. And this means that number r is the maximum number of linearly independent rows.

2.3.3. The sum and the intersection of subspaces. Suppose L_1 and L_2 are two arbitrary subspaces of one and the same linear space R .

The collection of all elements x of the space R belonging simultaneously to L_1 and L_2 forms a subspace of the space R^{**} called the intersection of the subspaces L_1 and L_2 .

* In particular, this statement yields a rather nontrivial theorem stating that in any matrix the maximum number of linearly independent rows coincides with the maximum number of linearly independent columns.

** Since the elements of this collection meet requirements 1° and 2° formulated at the beginning of 2.3.1.

The collection of all elements of the space R of the form $y+z$, where y is an element of the subspace L_1 and z is an element of the subspace L_2 , forms a subspace of the space R^* called the **sum** of the subspaces L_1 and L_2 .

Example. Suppose R is a linear space of all free vectors (in a three-dimensional space), L_1 is a subspace of all free vectors which are parallel to the plane Oxy and L_2 is a subspace of all free vectors which are parallel to the plane Oxz . Then the sum of the subspaces L_1 and L_2 is the whole space R^{**} and the intersection of the subspaces L_1 and L_2 is the set of all free vectors which are parallel to the axis Ox .

The following statement holds true.

Theorem 2.9. *The sum of the dimensions of the arbitrary subspaces L_1 and L_2 of the finite-dimensional linear space R is equal to the sum of the dimension of the intersection of these subspaces and the dimension of the sum of these subspaces.*

Proof. Let us designate the intersection of L_1 and L_2 as L_0 and the sum of L_1 and L_2 as $\hat{L}$. Considering L_0 to be R -dimensional, we choose in it a basis

$$e_1, e_2, \dots, e_k. \quad (2.11)$$

* See the preceding footnote.

** Indeed, any vector x of the space R is a linear combination $x = \alpha i + \beta j + \gamma k$ of the base vectors i, j, k , which are parallel to the axes Ox, Oy and Oz respectively, the vector $\alpha i + \beta j$ belonging to L_1 and the vector γk to L_2 .

Using the assertion proved in 2.3.1, we complete the basis (2.11) to form a basis

$$e_1, \dots, e_k, g_1, \dots, g_l \quad (2.12)$$

in the subspace L_1 , and to form a basis

$$e_1, \dots, e_k, f_1, \dots, f_m \quad (2.13)$$

in the subspace L_2 .

It is sufficient to prove that the elements

$$g_1, \dots, g_l, e_1, \dots, e_k, f_1, \dots, f_m \quad (2.14)$$

are the basis of the sum $\hat{L}$ of the subspaces L_1 and L_2 . To do that, it is, in turn, sufficient to prove that elements (2.14) are linearly independent and that any element x of the sum $\hat{L}$ is a certain linear combination of elements (2.14).

We shall first prove that elements (2.14) are linearly independent. Let us assume that a certain linear combination of elements (2.14) is a zero element, i.e. the equality

$$\alpha_1 g_1 + \dots + \alpha_l g_l + \beta_1 e_1 + \dots + \beta_k e_k + \gamma_1 f_1 + \dots + \gamma_m f_m = 0 \quad (2.15)$$

or

$$\alpha_1 g_1 + \dots + \alpha_l g_l + \beta_1 e_1 + \dots + \beta_k e_k = -\gamma_1 f_1 - \dots - \gamma_m f_m \quad (2.16)$$

holds true. Since the left-hand side of (2.16) is an element of L_1 and the right-hand side of (2.16) is an element of L_2 , it follows that both the left-hand and right-hand sides of (2.16) belong to the intersection L_0 of the subspaces L_1 and L_2 . Hence it follows, in particular, that the right-hand side of (2.16) is a certain linear combination of elements (2.11), i.e. there are numbers $\lambda_1, \dots, \lambda_R$ such that

* Since in this case the dimension of $\hat{L}$, equal to $i + k + m$, together with the dimension of L_0 , equal to k , is equal to the sum of the dimensions $k + l$ and $k + m$ of the subspaces L_1 and L_2 .

$$-\gamma_1 f_1 - \dots - \gamma_m f_m = \lambda_1 e_1 + \dots + \lambda_k e_k. \quad (2.17)$$

By virtue of linear independence of base elements (2.13), equality (2.17) is only possible when all the coefficients $\gamma_1, \dots, \gamma_m, \lambda_1, \dots, \lambda_k$ are equal to zero. But in this case we find from (2.15) that

$$\alpha_1 g_1 + \dots + \alpha_l g_l + \beta_1 e_1 + \dots + \beta_k e_k = 0. \quad (2.18)$$

Because of the linear independence of the base vectors (2.12), equality (2.18) is only possible when all the coefficient $\alpha_1, \dots, \alpha_l, \beta_1, \dots, \beta_k$ are equal to zero. Thus, we have established that equality (2.15) is possible only when all the coefficients $\alpha_1, \dots, \alpha_l, \beta_1, \dots, \beta_k, \gamma_1, \dots, \gamma_m$ are equal to zero, and this proves the linear independence of elements (2.14).

It remains to prove that any element x of the sum $\hat{L}$ is a certain linear combination of elements (2.14), but this follows immediately from the fact that the element x is (by the definition of $\hat{L}$) the sum of a certain element x_1 of the subspace L_1 , which is a linear combination of elements (2.12), and a certain element x_2 of the subspace L_2 , which is a linear combination of elements (2.13). We have proved the theorem.

Returning to the example considered prior to the formulation of Theorem 2.9, we note that the dimension of L_1 and that of L_2 is two, the dimension of their sum is three, and that of their intersection is unity.

2.3.4. Expansion of a linear space into direct sum of subspaces.

Suppose R_1 and R_2 are two subspaces of the linear n -dimensional space R .

Definition. We say that the space R is a **direct sum** of the subspaces R_1 and R_2 if every element x of the space R can be **uniquely** represented as the sum

$$x = x_1 + x_2 \quad (2.19)$$

of the element $\mathbf{x}_1$ of the subspace R_1 and the element $\mathbf{x}_2$ of the subspace R_2 .

The fact that R is a direct sum of R_1 and R_2 is symbolized as $R = R_1 \oplus R_2$.

The latter equality is usually called the **expansion of the space R into a direct sum of the subspaces R_1 and R_2** .

Thus, the space R of all free vectors (in a three-dimensional space) can be expanded into a direct sum of the subspace R_2 of all vectors parallel to the exist Oz .

Theorem 2.10. *For the n -dimensional space R to be a direct sum of the subspaces R_1 and R_2 it is sufficient that the intersection of R_1 and R_2 contain only a zero element and that the dimension of R be equal to the sum of the dimensions of the subspace R_1 and R_2 .*

Proof. Let us choose a basis $\mathbf{e}_1, \dots, \mathbf{e}_k$ in the subspace R_1 and a basis $\mathbf{g}_1, \dots, \mathbf{g}_l$ in the subspace R_2 . We shall prove that the union of these bases

$$\mathbf{e}_1, \dots, \mathbf{e}_k, \mathbf{g}_1, \dots, \mathbf{g}_l \quad (2.20)$$

is the basis of the whole space R . Since by the hypothesis the dimension n of the whole space R is equal to the sum $k+l$ of the dimensions of R_1 and R_2 , it is sufficient (by virtue of theorem 2.5) to prove the linear independence of elements (2.20).

Let us assume that a certain linear combination of elements (2.20) is zero elements, i.e. the equality

$$\alpha_1 \mathbf{e}_1 + \dots + \alpha_k \mathbf{e}_k + \beta_1 \mathbf{g}_1 + \dots + \beta_l \mathbf{g}_l = 0 \quad (2.21)$$

or

$$\alpha_1 \mathbf{e}_1 + \dots + \alpha_k \mathbf{e}_k = -\beta_1 \mathbf{g}_1 - \dots - \beta_l \mathbf{g}_l \quad (2.22)$$

holds true. Since the left-hand side of (2.22) is an element of R_1 and the right-hand side is an element of R_2 , and the intersection of R_1 and R_2 contains only a zero element, then the left-hand and the right-hand side of (2.25) are a zero element, and

this (because of linear independence of the elements of each of the bases $\mathbf{e}_1, \dots, \mathbf{e}_k$ and $\mathbf{g}_1, \dots, \mathbf{g}_l$) is possible only under the condition

$$\alpha_1 = \dots = \alpha_R = 0, \quad \beta_1 = \dots = \beta_l = 0. \quad (2.23)$$

Thus, we have found that equality (2.21) is possible only under the condition (2.23), and this proves the linear independence of elements (2.20) and the fact that elements (2.20) form a basis of the whole space R .

Suppose now that $\mathbf{x}$ is any element of R . Let us resolve it into components (2.20). We have $\mathbf{x} = \lambda_1 \mathbf{e}_1 + \dots + \lambda_k \mathbf{e}_k + \mu_1 \mathbf{g}_1 + \dots + \mu_l \mathbf{g}_l$ or $\mathbf{x} = \mathbf{x}_1 + \mathbf{x}_2$, where $\mathbf{x}_1 = \lambda_1 \mathbf{e}_1 + \dots + \lambda_k \mathbf{e}_k$ is an element of R_1 and $\mathbf{x}_2 = \mu_1 \mathbf{g}_1 + \dots + \mu_l \mathbf{g}_l$ is an element of R_2 .

It remains to prove that representation (2.19) is unique. Let us assume that besides (2.19) there holds another representation

$$\mathbf{x} = \mathbf{x}'_1 + \mathbf{x}'_2, \quad (2.24)$$

where $\mathbf{x}'_1$ is a element of R_1 and $\mathbf{x}'_2$ is an element of R_2 . Subtracting (2.24) from (2.19), we find that $0 = \mathbf{x}_1 - \mathbf{x}'_1 + \mathbf{x}_2 - \mathbf{x}'_2$, or $\mathbf{x}_1 - \mathbf{x}'_1 = \mathbf{x}_2 - \mathbf{x}'_2$. Since there is an element of R_1 on the left-hand side and since the intersection of R_1 and R_2 contains only a zero element, it follows from that equation that $\mathbf{x}_1 - \mathbf{x}'_1 = 0$, $\mathbf{x}_2 - \mathbf{x}'_2 = 0$, i.e. $\mathbf{x}'_1 = \mathbf{x}_1$, $\mathbf{x}'_2 = \mathbf{x}_2$.

And this is what we had to prove.

Remark. When the space R is not a straight line but an ordinary sum of the subspaces R_1 and R_2 , representation (2.19) of any element $\mathbf{x}$ of the space R is also valid but is not, in a general case, unique.

Suppose, for instance, that R is a three-dimensional space of all free vectors, R_1 is a subspace of all vectors which are parallel to the plane Oxy and R_2 is a subspace of all vectors which are parallel to the plane Oxz . We have found in 2.3.3 that R is the sum (but not a direct sum, of course) of the subspace R_1 and R_2 . Let us designate

2.4. Transformation of Coordinates in Transformation of the Basis of an n -Dimensional Linear Space

$$\left. \begin{aligned} \mathbf{e}'_1 &= \mathbf{a}_{11}\mathbf{e}_1 + \mathbf{a}_{12}\mathbf{e}_2 + \dots + \mathbf{a}_{1n}\mathbf{e}_n \\ \mathbf{e}'_2 &= \mathbf{a}_{21}\mathbf{e}_1 + \mathbf{a}_{22}\mathbf{e}_2 + \dots + \mathbf{a}_{2n}\mathbf{e}_n \\ &\vdots \\ \mathbf{e}'_n &= \mathbf{a}_{n1}\mathbf{e}_1 + \mathbf{a}_{n2}\mathbf{e}_2 + \dots + \mathbf{a}_{nn}\mathbf{e}_n. \end{aligned} \right\} \quad (2.25)$$
$$A = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{pmatrix}. \quad (2.26)$$

64

* In 1.2.7 we agreed to call a matrix of this kind non-singular.

Let us verify that *an inverse change from the second basis $e'_1, e'_2, \dots, e'_n$ to the first basis $e_1, e_2, \dots, e_n$ is carried out with the aid of the matrix which is an inverse of the matrix A .*

Recall that the matrix A which is an inverse of the matrix A was introduced in 1.2.7 and has the form

$$B = \begin{vmatrix} \frac{A_{11}}{\Delta} & \frac{A_{21}}{\Delta} & \dots & \frac{A_{n1}}{\Delta} \\ \frac{A_{12}}{\Delta} & \frac{A_{22}}{\Delta} & \dots & \frac{A_{n2}}{\Delta} \\ \dots & \dots & \dots & \dots \\ \frac{A_{1n}}{\Delta} & \frac{A_{2n}}{\Delta} & \dots & \frac{A_{nn}}{\Delta} \end{vmatrix}. \quad (2.27)$$

where Δ is the determinate of the matrix A and A_{ik} is a cofactor of the element a_{ik} of that determinant.

Let us multiply equation (2.25) by the respective cofactors $A_{1j}, A_{2j}, \dots, A_{nj}$ of the elements of the j th column of the determinant Δ and then add the equations together. As a result, we get (for any number j equal to $1, 2, \dots, n$)

$$e'_1 + e'_2 A_{2j} + \dots + e'_n A_{nj} = \sum_{i=1}^n e_i (a_i A_{1j} + a_{2i} A_{2j} + \dots + a_{ni} A_{nj}).$$

Taking into account that the sum of the products of the elements of the j th column by the respective cofactors of the elements of the j th column is equal to zero for $i \neq j$ and to the determinant Δ for $i = j^*$, we get from the last equation

$$e'_1 A_{1j} + e'_2 A_{2j} + \dots + e'_n A_{nj} = e_j \Delta,$$

* See property 4o in 1.2.4.

whence $e_j = \frac{A_{1j}}{\Delta} e'_1 + \frac{A_{2j}}{\Delta} e'_2 + \dots + \frac{A_{nj}}{\Delta} e'_n$ ($j = 1, 2, \dots, n$) or, in more detail,

(2.28)

abbreviated as A^{-1} .

that

$$\mathbf{x} = x'_1 \mathbf{e}'_1 + x'_2 \mathbf{e}'_2 + \dots + x'_n \mathbf{e}'_n = x_1 \mathbf{e}_1 + x_2 \mathbf{e}_2 + \dots + x_n \mathbf{e}_n.$$

by formulas (2.28), we obtain

$$\begin{aligned} \mathbf{x} &= x'_1 \mathbf{e}'_1 + x'_2 \mathbf{e}'_2 + \dots + x'_n \mathbf{e}'_n \\ &= x_1 \left(\frac{A_{11}}{\Delta} \mathbf{e}'_1 + \frac{A_{21}}{\Delta} \mathbf{e}'_2 + \dots + \frac{A_{n1}}{\Delta} \mathbf{e}'_n \right) \\ &\quad + x_2 \left(\frac{A_{12}}{\Delta} \mathbf{e}'_1 + \frac{A_{22}}{\Delta} \mathbf{e}'_2 + \dots + \frac{A_{n2}}{\Delta} \mathbf{e}'_n \right) \end{aligned}$$

$$+x_n\left(\frac{A_{1n}}{\Delta}\mathbf{e}'_1+\frac{A_{2n}}{\Delta}\mathbf{e}'_2+\dots+\frac{A_{nn}}{\Delta}\mathbf{e}'_n\right).$$

By virtue of the uniqueness of the resolution into the components $e'_1, e'_2, \dots, e'_n$, the last equality immediately yields formulas for transition from the coordinates $(x_1, x_2, \dots, x_n)$ with respect to the first basis to the coordinates $(x'_1, x'_2, \dots, x'_n)$ with respect to the second basis:

$$\left. \begin{aligned} x_1' &= \frac{A_{11}}{\Delta}x_1 + \frac{A_{12}}{\Delta}x_2 + \dots + \frac{A_{1n}}{\Delta}x_n, \\ x_2' &= \frac{A_{21}}{\Delta}x_1 + \frac{A_{22}}{\Delta}x_2 + \dots + \frac{A_{2n}}{\Delta}x_n, \\ &\vdots \\ x_n' &= \frac{A_{n1}}{\Delta}x_1 + \frac{A_{n2}}{\Delta}x_2 + \dots + \frac{A_{nn}}{\Delta}x_n \end{aligned} \right\} \quad (2.29)$$

Formulas (2.29) show that the transition from the coordinates $(x_1, x_2, \dots, x_n)$ to the coordinates $(x'_1, x'_2, \dots, x'_n)$ is carried out by means of the matrix

$$C = \begin{vmatrix} \frac{A_{11}}{\Delta} & \frac{A_{12}}{\Delta} & \dots & \frac{A_{1n}}{\Delta} \\ \frac{A_{21}}{\Delta} & \frac{A_{22}}{\Delta} & \dots & \frac{A_{2n}}{\Delta} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{A_{n1}}{\Delta} & \frac{A_{n2}}{\Delta} & \dots & \frac{A_{nn}}{\Delta} \end{vmatrix}$$

transposed to the inverse matrix (2.27)

We arrive at the following conclusion: *if a transition from the first basis to the second is carried out by means of the non-singular matrix A , then the transition from the coordinates of that element with respect to the second basis is carried out by means of the matrix $(A^{-1})'$ transposed to the inverse matrix A^{-1} .*

CHAPTER – 3

SYSTEMS OF LINEAR EQUATIONS

From the elementary course of linear algebra and from analytic geometry the reader is sure to be familiar with a system of two linear equations in two unknown and with systems of two or three linear equations in three unknowns*. In this chapter we study systems of arbitrary numbers m of linear equations in the arbitrary number n of unknowns.

We first establish the necessary and sufficient condition for the existence of at least one solution (or, as they say, the **consistency**) of such a system and then seek the whole collection of its solutions.

In 4.4. we consider the case of an approximate specification of all the coefficients of the system and its constant terms which is important for applications. For that case, we present **Tikhonov's regularization method** which makes it possible to find the so-called **normal** (i.e. the closest of the origin) solution of the indicated system with an accuracy corresponding to the accuracy of specification of coefficient and constant terms.

Chapter 6 gives an idea of the numerical (iteration) methods of solving systems of linear equations.

3.1. The Condition for the Consistency of a Linear System

3.1.1. The concept of a system of linear equations and its solutions.

In a general case a system of m linear equations in n unknowns (or briefly, a **linear system**) has the form

$$\left. \begin{aligned} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n &= b_1, \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n &= b_2, \\ . &. \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n &= b_m. \end{aligned} \right\}. \quad (3.1)$$

The designation $x_1, x_2, \dots, x_n$, are used for the unknowns to be defined (their number n is not assumed to be obligatory equal to the number of equations m), the quantities $a_{11}, a_{12}, \dots, a_{mn}$ called the **coefficients of the system** and the quantities $b_1, b_2, \dots, b_m$, called the **free** or **constant terms**, are assumed to be known. Every coefficient a_{ij} of the system has two indices, the first of which i indicates the number of the equation and the second j indicates the numbers of the unknown in which this coefficient is.

System (3.1) is called **homogenous** if all its constant terms $b_1, b_2, \dots, b_m$ are equal to zero.

If at least one of the constant terms $b_1, b_2, \dots, b_m$ is different from zero, then system (3.1) is called **nonhomogeneous**.

System (3.1) said to be **quadratic** if the number m of its equations is equal to the number n of its unknowns.

The **solution** of system (3.1) is such a collection n of the numbers $c_1, c_2, \dots, c_n$, which, being substituted into system (3.1) for the unknowns $x_1, x_2, \dots, x_n$, turns all the equations of the system into identities.

Not every system of the form (3.1) has solutions. Thus, the systems of linear equations

$$\left. \begin{array}{l} x_1 + x_2 = 1, \\ x_1 + x_2 = 2 \end{array} \right\}$$

obviously has no solutions (since if a solution of the system existed, then, upon its substitution into the system, the left-hand sides of the two equations would contain identical numbers and we would have $1 = 2$).

The system of equations of the form (3.1) is called **consistent** if it has at least one Solutions. The consistent system of the form (3.1) can have either one solution or several solutions.

Two solution $c_1^{(1)}, c_2^{(1)}, \dots, c_n^{(1)}$ and $c_1^{(2)}, c_2^{(2)}, \dots, c_n^{(2)}$ of the consistent system of the form (3.1) are said to be distinct if at least one of the equalities $c_1^{(1)} = c_1^{(2)} = c_2^{(1)} = c_2^{(2)}, \dots, c_n^{(1)} = c_n^{(2)}$ is violated.

The consistent system of the form (3.1) is called **determinate** if it has a unique solution.

The consistent system of the form (3.1) is called **indeterminate** if it has at least two distinct solutions.

It is convenient to write the linear system (3.1) in matrix form. For that purpose we use the notion of the product of two matrices (such that the number of the columns of the first matrix is equal to the number of the rows of the second matrix) introduced in 1.1.2. We take the following two matrices as the matrices being multiplied: the matrix

$$A \begin{vmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{vmatrix}, \quad (3.2)$$

containing m rows and n columns and formed from the coefficients in the unknowns (in what follows we call such a matrix a **coefficient matrix** of system (3.11), and the matrix X consisting of n rows and 1 column, i.e. a column of the form)

$$X = \begin{vmatrix} x_1 \\ x_2 \\ \dots \\ x_n \end{vmatrix}. \quad (3.3)$$

According to the rule of multiplication of two matrices, the product AX of matrix (3.2) by matrix (3.3) is a matrix consisting of m rows and 1 column, i.e. a column of the form

$$\begin{pmatrix} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n \\ \dots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n \end{pmatrix}. \quad (3.4)$$

The system of equalities (3.1) means that the column (3.4) coincides with the column

$$B = \begin{pmatrix} b_1 \\ b_2 \\ \dots \\ b_m \end{pmatrix}. \quad (3.5)$$

Thus, in matrix notation, system (3.1) can be replaced by one matrix equation

$$AX = B, \quad (3.6)$$

which is equivalent of it and in which the matrices A , X and B are defined by relations (3.2), (3.3) and (3.5). The solution of the matrix equation (3.6) consists in seeking a column (3.3) which turns equation (3.6) into identity when matrix (3.2) and column (3.5) of the right-hand sides are given.

In this and in the following section we shall elucidate the following three question concerning system (3.1):

- (1) the way to find out whether system (3.1) is consistent,
- (2) the way to find out whether system (3.1) (when it is consistent) is determinate,
- (3) the way of seeking the unique solution of the consistent system (3.1) (when it is determinate) and seeking all its solutions (when it is indeterminate).

3.1.2 non-trivial consistency of a homogenous system. We begin with considering a homogenous linear system of the form (3.1), i.e. the system

$$\left. \begin{aligned} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n &= 0, \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n &= 0, \\ \dots &\dots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n &= 0. \end{aligned} \right\} \quad (3.7)$$

We note that this system is *always consistent* since it always has the so-called **trivial** (or **zero**) solution $x_1 = x_2 = \dots = x_n = 0^*$. (*Indeed, substituting in system (3.7) zeros for all the unknowns $x_1, x_2, \dots, x_n$, we turn all the equations of the system into identities.)

A question arises as to the *condition under which the homogeneous system (3.7) has other solutions in addition to the trivial solution* (i.e. is **non-trivial consistent**).

This question is simple to solve. Note that the existence of a nontrivial solution of system (3.7) is equivalent to a linear dependence of the columns of matrix (3.2) (since the linear dependence of the columns of matrix (3.2) means that there are numbers $x_1, x_2, \dots, x_n$, *not all equal to zero*, such that equalities (3.7) hold true).

But by Theorem 1.6 on the base minor, the linear dependence of the columns of matrix (3.2) exists if and only if **not all** the columns of that matrix are basic, i.e. if and only if the order r of the base minor of matrix (3.2) is smaller than the number n of its columns.

We arrive at the following theorem.

Theorem 3.1. *The homogeneous system (3.7) has non-trivial solutions if and only if the rank r of matrix (3.2) is smaller than the number n of its columns.*

Corollary. *A quadratic homogeneous system has non-trivial solutions if and only if the determinant formed from the coefficients in the unknowns is equal to zero.*

Indeed, in the case of the quadratic homogenous system (2.7) i.e. for $m = n$, the rank r of matrix (3.2) is smaller than number $m = n$, if and only if the determinant of that matrix is equal to zero.

3.1.3. The condition of the consistency of a general linear system. Let us now establish the necessary and sufficient condition for the consistency of a general (inhomogeneous in a general case) system of the form (3.1). This system is associated with two matrices: a matrix A defined by relation (3.2) usually called a **system** or *coefficient matrix* of (3.1) (it is formed from the coefficients in the unknowns), and a matrix

$$A_1 = \left\| \begin{array}{cccc} a_{11} & a_{12} & \dots & a_{1n}b_1 \\ a_{21} & a_{22} & \dots & a_{2n}b_2 \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn}b_m \end{array} \right\|, \quad (3.8)$$

which is usually called an **augmented matrix of system** (3.1) (it results from the coefficient matrix by adding column (3.5) of constant terms to that matrix).

The following fundamental theorem holds true.

Theorem 3.2 (the Kronecker-Capelli theorem). *For the linear system (3.1) to be consistent, it is necessary and sufficient that the rank of the augmented matrix of that system be equal to that of the coefficient matrix.*

Proof. (1) Necessity. Suppose system (3.1) is consistent, i.e. there are numbers $c_1, c_2, \dots, c_n$ such that the equalities

$$\left. \begin{array}{l} a_{11}c_1 + a_{12}c_2 + ... + a_{1n}c_n = b_1, \\ a_{21}c_1 + a_{22}c_2 + ... + a_{2n}c_n = b_2, \\ . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \quad . \\ a_{m1}c_1 + a_{m2}c_2 + ... + a_{mn}c_n = b_m \end{array} \right\} \quad (3.9)$$

hold true. Let us designate as r the rank of the **matrix of coefficients** (3.1) and consider the linear closure L of the r basic columns of that matrix. By virtue of Theorem 1.6 on the base minor; any column of the matrix of coefficients belongs to the indicated linear closure L . In other words, any column of the augmented matrix (3.8), except for its last column, belongs to the indicated linear closure L .

It follows from equations (3.9) that the last column of the augmented matrix (3.8) also belongs to the linear closure L (since by virtue of equations (3.9), this last column can be linearly expressed in terms of all the columns of the matrix of coefficients and can, therefore, be linearly expressed in terms of its base columns).

Thus, *all the columns of the augmented matrix (3.8) belong to the indicated linear closure L* . We have established, in 2.3.2, that the dimension of the indicated linear closure L is equal to r .

This means that any $r + 1$ columns of the augmented matrix (3.8) are linearly dependent, i.e. the rank of the augmented matrix (equal maximum number of

(2) Sufficiency. Suppose the ranks of the system matrix and the augmented matrix coincide. Then r base columns of the system matrix are also the base columns of the augmented matrix (3.8). By Theorem 1.6, on the base minor, the last column of the augmented matrix (3.8) is a certain linear combination of the indicated r base columns. Hence the last column of the augmented matrix (3.8) is also a linear combination of all the columns of the system matrix (3.2), i.e. there are numbers $c_1, c_2, \dots, c_n$, such that equalities (3.9) hold true. The last equalities signify that the numbers $c_1, c_2, \dots, c_n$ are a solution of system (3.1), i.e. that system is consistent. This completes the proof of the theorem.

The Kronecker-Capelli theorem establishes the necessary and sufficient condition for a linear system to be consistent but does not present any means of finding the solutions of that system. We shall occupy ourselves here with seeking the solutions of the linear system (3.1). We shall first consider the simplest case of a quadratic system of linear equations with the determinants of the system matrix different from zero and pass to seeking the collection of the solutions of the general linear system of the form (3.1).

$$\left. \begin{array}{l} a_{11}x_1 + a_{12}x_2 + ... + a_{1n}x_n = b_1, \\ a_{21}x_1 + a_{22}x_2 + ... + a_{2n}x_n = b_2, \\ .\quad.\quad.\quad.\quad.\quad.\quad.\quad.\quad.\quad.\quad. \\ a_{n1}x_1 + a_{n2}x_2 + ... + a_{nn}x_n = b_n \end{array} \right\} \quad (3.10)$$
$$A = \begin{vmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{vmatrix} \quad (3-11)$$

We shall prove that such a system possesses a solution and that the solution is unique. Let us now find that solution. We shall first prove that system (3.10) can have only one solution (i.e. we shall prove the uniqueness of the solution of system (3.10) under the assumption that it exists).

We assume that there are any n numbers $x_1, x_2, \dots, x_n$ such that being substituted into system (3.10), they turn all the equations of the system into identities (i.e. there is a certain $x_1, x_2, \dots, x_n$ of system (3.10)). Then, multiplying identities (3.10) by the respective cofactors $A_{1j}, A_{2j}, \dots, A_{nj}$, of the elements of the j th column of the determinant Δ of matrix (3.11) and then adding the resulting identities together, we find (for any number j equal to $1, 2, \dots, n$)

$$\sum_{i=1}^n x_i (a_{1j} + a_{2j}A_{2j} + \dots + a_{nj}A_{nj} = b_1A_{1j} + b_2A_{2j} + \dots + b_nA_{nj}).$$

Taking into account that the sum of the products of the elements of the i th column by the corresponding cofactors of the elements of the j th column is equal to zero at $i \neq j$ and is equal to the determinant Δ of matrix (3.11) at $i = j$, we get from the last equation.

$$x_j \Delta = b_1A_{1j} + b_2A_{2j} + \dots + b_nA_{nj}. \quad (3.12)$$

We shall use the designation $\Delta_j(b_i)$ (or, briefly, Δ_i) for the determinant resulting from the determinant Δ of the coefficient matrix of system (3.11) by replacing its j th column by the column of the constant terms $b_1, b_2, \dots, b_n$ (leaving all the other columns of Δ intact).

Note that the right-hand side of (3.12) contains the determinant $\Delta_j(b_i)^{**}$, (** To ascertain this fact, it is sufficient to write the expansion of the determinant $\Delta_j(b_i)$ according to the elements of the i th column.) and the equation assumes the form

$$x_j \Delta = \Delta_i \quad (i = 1, 2, \dots, n), \quad (3.13)$$

The determinant Δ of matrix (3.11) being nonzero, equations (3.13) are equivalent to the relations

$$x_j = \frac{\Delta_j}{\Delta} \quad (j = 1, 2, \dots, n). \quad (3.14)$$

We have thus proved that *if the solution $x_1, x_2, \dots, x_n$ of system (3.10) with the determinant Δ of the system matrix (3.11) different from zero exists, then that solution is uniquely defined by formulas (3.14).*

Formulas (3.14) are known as **Cramer's rule**.

It should be emphasized, once more that Cramer's formulas have been obtained under the assumption of the existence of a solution and prove its uniqueness.

It remains to prove the existence of a solution of system (3.10). To do that, it is sufficient, by virtue of the Kronecker-Capelli theorem, to prove that the rank of the system matrix (3.11) is equal to that of the augmented matrix

$$A_1 = \left\| \begin{array}{cccc} a_{11} & a_{12} & \dots & a_{1n}b_1 \\ a_{21} & a_{22} & \dots & a_{2n}b_2 \\ \cdot & \cdot & \cdot & \cdot \\ a_{n1} & a_{n2} & \dots & a_{nn}b_n \end{array} \right\|, \quad (3.15)$$

but this is obvious since by the relation $\Delta \neq 0$ the rank of the system matrix is n and the rank of the augmented matrix (3.15) containing n rows cannot exceed the number n and is, therefore, equal to the rank of the system matrix.

We have thus completely proved *that the quadratic system of linear equations (3.10) with the determinant of the system matrix different from zero has a solution defined by Cramer's rule (3.14) and that the solution is unique.*

It is even easier to verify the assertion we have proved by matrix method. To do that, we replace (as in 1.1) system (3.10) by the matrix equation equivalent to it:

$$AX = B, \quad (3.16)$$

where A is the matrix of the coefficients of (3.11) and X and B are the columns

$$X = \left\| \begin{array}{c} x_1 \\ x_2 \\ \cdot \\ x_n \end{array} \right\|, \quad B = \left\| \begin{array}{c} b_1 \\ b_2 \\ \cdot \\ b_n \end{array} \right\|,$$

the first of which must be found and the second is given.

Since the determinant Δ of the matrix A is nonzero, there is an inverse matrix A^{-1} (see 1.2.7).

Let us assume that there is a solution of system (3.10), i.e. there is a column X turning the matrix equation (3.16) into identity. Multiplying the indicated identity from the left by the inverse matrix A^{-1} , we get

$$A^{-1} (AX) = A^{-1} B. \quad (3.17)$$

we must now take into account that by virtue of the associative property of the product of three matrices (see 1.1.2) and according to the relation $A^{-1} A = I$, where I is an identity matrix (see 1.2.7), $A^{-1} (AX) = (A^{-1} A) X = IX = X$, so that we find from (3.17) that

$$X = A^{-1} B. \quad (3.18)$$

Writing out equation (3.18) and taking into account the form of the inverse matrix, we obtain Cramer's formulas for the elements of the column X .

Thus we have proved that if a solution of the matrix equation (3.16) exists, then it is uniquely defined by relation (3.18) equivalent to Cramer's formulas.

It is easy to verify that the column X defined by relation (3.18) is, indeed, a solution of matrix equation (3.16), i.e. being substituted in that equation, it turns it into identity. In fact, if the column X is defined by equation (3.18), then $AX = A (A^{-1} B) = AA^{-1} B = IB = B$.

Thus, *if the determinant Δ of the matrix A is nonzero (i.e. if the matrix is non-singular), then there is a solution of matrix equation (3.16), defined by relation (3.18) which is equivalent to Cramer's formulas, and that solution is unique.*

Example. Let us find the solution of the quadratic. system of linear equations

$$\left. \begin{aligned} x_1 + 2x_2 + 3x_3 + 4x_4 &= 30, \\ -x_1 + 2x_2 - 3x_3 + 4x_4 &= 10, \\ x_2 - x_3 + x_4 &= 3, \\ x_1 + x_2 + x_3 + x_4 &= 10, \end{aligned} \right\}$$

with a nonzero determinant of the system matrix

$$\Delta = \begin{vmatrix} 1 & 2 & 3 & 4 \\ -1 & 2 & -3 & 4 \\ 0 & 1 & -1 & 1 \\ 1 & 1 & 1 & 1 \end{vmatrix} = -4.$$

Since

$$\Delta_1 = \begin{vmatrix} 30 & 2 & 3 & 4 \\ 10 & 2 & -3 & 4 \\ 3 & 1 & -1 & 1 \\ 10 & 1 & 1 & 1 \end{vmatrix} = -4, \quad \Delta_2 = \begin{vmatrix} 1 & 30 & 3 & 4 \\ -1 & 10 & -3 & 4 \\ 0 & 3 & -1 & 1 \\ 1 & 10 & 1 & 1 \end{vmatrix} = -8.$$

$$\Delta_3 = \begin{vmatrix} 1 & 2 & 30 & 4 \\ -1 & 2 & 10 & 4 \\ 0 & 1 & 3 & 1 \\ 1 & 1 & 10 & 1 \end{vmatrix} = -4, \quad \Delta_4 = \begin{vmatrix} 1 & 2 & 3 & 30 \\ -1 & 2 & -3 & 10 \\ 0 & 1 & -1 & 3 \\ 1 & 1 & 1 & 10 \end{vmatrix} = -16.$$

the only solution of the system in question has the form $x_1 = 1, x_2 = 2, x_3 = 3, x_4 = 4$ by virtue of Cramer's rule.

The main significance of Cramer's rule consists in the fact that it gives an explicit expression for the solution of a quadratic system of linear equations (with a nonzero determinant) in terms of the coefficients of the equations and constant terms. The practical application of Cramer's rule is connected with cumbersome calculations (($n + 1$) determinants of order n must be calculated to solve a system of n equations in n unknowns). It should be added that if the coefficients of the equations and the constant terms are only approximate values of some physical quantities being measured or are rounded off in the process of calculation, then the use of Cramer's rule may lead to major errors and is impractical in a number of cases.

Tikhonov's regularization method, which will be presented in 4.4, makes it possible to find solutions of a linear system with an accuracy corresponding to the accuracy of the specification of the matrix of coefficients of the equations and of the column of constant terms, and in Chapter 6 we shall give an idea of the so-called iteration methods of solving linear systems which make it possible to find solutions for systems of that kind by successive approximations of the unknowns.

3.2.2. Seeking all the solutions of a general linear system. Let us now consider a general system of m linear equations in n unknowns (3.1). We assume that this system is consistent and that the rank of its system and augmented matrices is equal to r . Without losing generality, we can assume that the base minor of system matrix (3.2) is in the top left-hand corner of that matrix (the general case can be reduced to this case by transposing the equations and the unknowns in system (3.1)).

In terms of system (3.1), this means that each of the equations of this system, beginning with the $(r + 1)$ th equation, is a linear combination (i.e. the consequence) of the first r equations of that system (i.e. *every solution of the first r equations of system (3.1) turns into identities all the other equations of the system*).

[illegible]

79

$c_{r+1}, \dots, c_n$ there is a unique collection of r numbers $c_1, c_2, \dots, c_r$ turning all the equations of system (3.19) into identities and defined by Cramer's rule.

To put down this unique solution, we shall agree to use the designation $M_j(d_i)$ for the determinant obtained from the base minor M of matrix (3.2) by replacing its j th column by the column of numbers $d_1, d_2, \dots, d_i, \dots, d_r$ (leaving all the other columns of M unchanged). Then, writing the solution of system (3.19) with the aid of Cramer's formulas and using the linear property of the determinant, we get,

$$\begin{aligned} c_j &= \frac{1}{M} M_j(b_i - a_{i(r+1)}c_{r+1} - \dots - a_{in}c_n) \\ &= \frac{1}{M} [M_j(b_i) - c_{r+1} M_j(a_{i(r+1)}) - \dots - c_n M_j(a_{in})] \\ &\quad (j = 1, 2, \dots, r). \end{aligned}$$

Formulas (3.20) express the values of the unknowns $x_i = c_i$ ($i=1, 2, \dots, r$) in terms of the coefficients in the unknowns, the constant terms and arbitrary parameters $c_{r+1}, \dots, c_n$.

Let us prove that *formulas (3.20) contain any solution of system (3.1)*. Indeed, suppose $c_1^0, c_2^0, \dots, c_r^0, c_{r+1}^0, \dots, c_n^0$ is an arbitrary solution of the indicated system. Then it is also a solution of system (3.19). But the quantities $c_1^0, c_2^0, \dots, c_r^0$ appearing in system (3.19) can be uniquely expressed in terms of the quantities $c_{r+1}^0, \dots, c_n^0$ by Cramer's formulas (3.20). Thus, for $c_{r+1} = c_{r+1}^0, \dots, c_n = c_n^0$ formulas (3.20) give the solution $c_1^0, c_2^0, \dots, c_r^0, c_{r+1}^0, \dots, c_n^0$ being considered.

Remark. If the rank r of the system matrix and the augmented matrix of system (3.1) is equal to the number n of the unknowns, then in this case as well relations (3.20) pass into the formulas

$$c_j = \frac{M_j(b_i)}{M} \quad (j = 1, 2, \dots, n).$$

defining the **unique** solution of system (3.1). Thus, *system (3.1.) possesses a unique solution (Le. it is determinate) under the condition that the rank r of the system matrix and the augmented matrix is equal to the number of the unknowns n (and is less than the number m of equations or is equal to it).*

Example. Let us find all the solutions of the linear system

$$\left. \begin{aligned} x_1 - x_2 + x_3 - x_4 &= 4, \\ x_1 + x_2 + 2x_3 + 3x_4 &= 8, \\ 2x_1 + 4x_2 + 5x_3 + 10x_4 &= 20, \\ 2x_1 - 4x_2 + x_3 - 6x_4 &= 4. \end{aligned} \right\} \quad (3.21)$$

It is easy to ascertain that the rank of both the coefficient matrix and the augmented matrix of this system is equal to 2 (i.e. the system is consistent) and it may be considered that the base minor M is in the top left corner of the system matrix, i.e.

$M = \begin{vmatrix} 1 & -1 \\ 1 & 1 \end{vmatrix} = 2$. But then, deleting the last two equations and assigning arbitrary

values to c_3 and c_4 , we get a system

$$\left. \begin{aligned} x_1 - x_2 &= 4 - c_3 + c_4, \\ x_1 + x_2 &= 8 - 2c_3 - 3c_4, \end{aligned} \right\}$$

from which, by virtue of Cramer's rule, we get the values

$$x_1 - c_1 = 6 - \frac{3}{2}c_3 - c_4, \quad x_2 = c_2 \quad 2 - \frac{1}{2}c_3 - 2c_4. \quad (3.22)$$

Thus, four numbers.

$$\left(6 - \frac{3}{2}c_3 - c_4, 2 - \frac{1}{2}c_3 - 2c_4, c_3, c_4 \right) \quad (3.23)$$

form a solution of system (3.21) for arbitrary values of c_3 and c_4 , the row (3.23) containing all the solutions of this system.

3.2.3. The properties of the collection of solutions of a homogeneous system. Let us now discuss a homogeneous system of m linear equations in n unknowns (3.7)

assuming, as before; that matrix (3.2) has a rank equal to r and that the base minor M is in the top left corner of that matrix.

Since in this case all b_i are equal to zero, we get the following formulas instead of (3.20):

$$c_1 = -\frac{1}{M} \left[c_{r+1} M_j(a_{i(r+1)}) + \dots + c_n M_j(a_{in}) \right] (j=1, 2, \dots, r) \quad (3.24)$$

which express the values of the unknowns $x_j = c_j$ ($j=1, 2, \dots, r$) in terms of the coefficients in the unknowns and the arbitrary values $c_{r+1}, \dots, c_n$. Because of what was proved in 3.2.2, *formulas (3.24) contain any solution of homogeneous system (3.7).*

Let us now make sure that the *collection of all solutions of homogeneous system (3.7) forms a linear space.*

Suppose $X_1 = (x_1^{(1)}, \dots, x_n^{(1)})$ and $X_2 = (x_1^{(2)}, \dots, x_n^{(2)})$ are two arbitrary solutions of homogeneous system (3.7) and λ is any real number. Since every solution of homogeneous system (3.7) is an element of the linear space A^n of all ordered collections of n numbers, it is sufficient to prove that each of the two collections

$$X_1 + X_2 = (x_1^{(1)} + x_1^{(2)}, \dots, x_n^{(1)} + x_n^{(2)}) \text{ and } \lambda X_1 = (\lambda x_1^{(1)}, \dots, \lambda x_n^{(1)})$$

is also a solution of homogeneous system (3.7).

Let us consider any equation of (3.7), say, the i th equation, and substitute the elements of the indicated collections for the unknowns in this equation. Taking into account that X_1 and X_2 are solutions of the homogeneous system, we get

$$\sum_{j=1}^n a_{ij} [x_1^{(1)} + x_1^{(2)}] = \sum_{j=1}^n a_{ij} x_j^{(1)} + \sum_{j=1}^n a_{ij} x_j^{(2)} = 0,$$

$$\sum_{j=1}^n a_{ij} [\lambda x_j^{(1)}] = \lambda \sum_{j=1}^n a_{ij} x_j^{(1)} = 0.$$

and this means that the collections $X_1 + X_2$ and λX_1 are solutions of homogeneous system (3.7).

Thus, the collection of all solutions of homogeneous system (3.7) forms a linear space which we symbolize as R .

Next we find the dimension of that space R and construct a basis in it.

We shall prove that under the assumption that the rank of the matrix of homogeneous system (3.7) is equal to r , the linear space R of all solutions of homogeneous system (3.7) is isomorphic to the linear space A^{n-r} of all the ordered collections of $(n-r)$ numbers.

Let us put every solution $(c_1, \dots, c_r, c_{r+1}, \dots, c_n)$ of homogeneous system (3.7) into correspondence with the element $(c_{r+1}, \dots, c_n)$ of the space A^{n-r} . Since the numbers $c_{r+1}, \dots, c_n$ can be chosen arbitrarily and in each case the solution of system (3.7) is defined uniquely with the aid of formulas (3.24), the correspondence we have established is **one-to-one**. Next we note that if the elements $(c_{r+1}^{(1)}, \dots, c_n^{(1)})$ and $(c_{r+1}^{(2)}, \dots, c_n^{(2)})$ of the space A^{n-r} correspond to the elements $(c_1^{(1)}, \dots, c_r^{(1)}, c_{r+1}^{(1)}, \dots, c_n^{(1)})$ and $(c_1^{(2)}, \dots, c_r^{(2)}, c_{r+1}^{(2)}, \dots, c_n^{(2)})$ of the space R , then it follows immediately from formulas (3.24) that the element $(c_{r+1}^{(1)} + c_{r+1}^{(2)}, \dots, c_n^{(1)} + c_n^{(2)})$ is associated with the element $(c_1^{(1)} + c_1^{(2)}, \dots, c_n^{(1)} + c_n^{(2)}, c_{r+1}^{(1)} + c_{r+1}^{(2)}, \dots, c_n^{(1)} + c_n^{(2)})$ and the element $(\lambda c_{r+1}^{(1)}, \dots, \lambda c_n^{(1)})$ is associated with the element $(\lambda c_1^{(1)}, \dots, \lambda c_r^{(1)}, \lambda c_{r+1}^{(1)}, \dots, \lambda c_n^{(1)})$ for any real λ . We have thus proved that the correspondence we have established is isomorphic.

Thus, *the linear space R of all solutions of homogeneous system (3.7) in n unknowns and with the rank of the system matrix equal to r is isomorphic to the space A^{n-r} and, hence, has a dimension $n-r$.*

*Any collection of $(n-r)$ linearly independent solutions of homogeneous system (3.7) form (by virtue of Theorem 2.5) a basis in the space R of all solutions and is known as a **fundamental collection of solutions** of homogeneous system (3.7).*

We must isolate the fundamental collection of solutions of system (3.7) corresponding to the simplest basis $e_1 = (1, 0, 0, \dots, 0)$, $e_2 = (0, 1, 0, \dots, 0), \dots, e_{n-r} = (0, 0, 0, \dots, 1)$ of the space A^{n-r} and known as a **normal fundamental collection of solutions** of homogeneous system (3.7).

$$\left. \begin{aligned} X_1 &= \left(-\frac{M_1(a_{i(r+1)})}{M}, \dots, -\frac{M_r(a_{i(r+1)})}{M}, 1, 0, \dots, 0 \right) \\ X_2 &= \left(-\frac{M_1(a_{i(r+2)})}{M}, \dots, -\frac{M_r(a_{i(r+2)})}{M}, 0, 1, \dots, 0 \right) \\ &\vdots \\ X_{n-r} &= \left(-\frac{M_1(a_{in})}{M}, \dots, -\frac{M_r(a_{in})}{M}, 0, 0, \dots, 1 \right) \end{aligned} \right\} \quad (3.25)$$
$$X = C_1 X_1 + C_2 X_2 + \dots + C_{n-r}, \quad (3.26)$$

Example. Let us consider a homogeneous system of equations

$$\left. \begin{aligned} x_1 - x_2 + x_3 - x_4 &= 0, \\ x_1 + x_2 + 2x_3 + 3x_4 &= 0, \\ 2x_1 + 4x_2 + 5x_3 + 10x_4 &= 0, \\ 2x_1 - 4x_2 + x_3 - 6x_4 &= 0. \end{aligned} \right\} \quad (3.27)$$

corresponding to nonhomogeneous system (3.21) discussed in the example at the end of 3.2.2. We have found there that the rank r of the matrix of that system is two and have taken the minor positioned at the top left corner of the indicated matrix as a base minor.

Repeating the arguments presented at the end of 3.2.2, we get, instead of formulas (3.22), relations

$$c_1 = -\frac{3}{2}c_3 - c_4, \quad c_2 = -\frac{1}{2}c_3 - 2c_4$$

valid for the arbitrarily chosen c_3 and c_4 . With the aid of these relations (setting first $c_3 = 1, c_4 = 0$ and then $c_3 = 0, c_4 = 1$), we obtain a normal fundamental collection of two solutions of system (3.27):

$$X_1 = \left(-\frac{3}{2}, -\frac{1}{2}, 1, 0\right), \quad X_2 = (-1, -2, 0, 1)$$

The general solution of homogeneous system (3.27) has the form

$$X = C_1 \left(-\frac{3}{2}, -\frac{1}{2}, 1, 0\right) + C_2 (-1, -2, 0, 1), \quad (3.28)$$

where C_1 and C_2 are arbitrary constants.

Now we can establish a connection between the solutions of nonhomogeneous linear system (3.1) and the corresponding homogeneous system (3.7). Let us prove the following **two assertions**.

1°. *The sum of any solution of nonhomogeneous system (3.1) and any solution of the corresponding homogeneous system (3.7) is a solution of system (3.1).*

Indeed, if $c_1, \dots, c_n$ is a solution of system (3.1) and $d_1, \dots, d_n$ is a solution of the corresponding homogeneous system (3.7), then substituting the numbers $c_1 + d_1, \dots, c_n + d_n$ for the unknowns in any (say, i th) equation of system (3.1), we get

$$\sum_{j=1}^n a_{ij}(c_j + d_j) = \sum_{j=1}^n a_{ij}c_j + \sum_{j=1}^n a_{ij}d_j = b_j + 0 = b_j.$$

And that is what we set out to prove.

2°. *The difference between two arbitrary solutions of non-homogeneous system (3.1) is a solution of the corresponding homogeneous system (3.7).*

Indeed, if $c'_1, \dots, c'_n$ and $c''_1, \dots, c''_n$ are two arbitrary solutions of system (3.1), then, substituting the numbers $c'_1 - c''_1, \dots, c'_n - c''_n$, for the unknowns in any (say, i th) equation of system (3.7), we get

$$\sum_{j=1}^n a_{ij}(c'_j - c''_j) = \sum_{j=1}^n a_{ij}c'_j - \sum_{j=1}^n a_{ij}c''_j = b_i - b_i = 0.$$

And that is what we wished to prove.

It follows from the assertion we have proved that *having found one solution of nonhomogeneous system (3.1) and adding it to every solution of the corresponding homogeneous system (3.7), we can get all the solutions of nonhomogeneous system (3.1). In other words, the sum of a particular solution of the non-homogeneous system (3.1) and the general solution of the corresponding homogeneous system (3.7) yields a general solution of nonhomogeneous system (3.1).*

It is natural to take as a particular solution of nonhomogeneous system (3.1) its solution

$$X_0 = \left(\frac{M_1(b_i)}{M}, \dots, -\frac{M_r(b_i)}{M}, 0, 0, \dots, 0 \right), \quad (3.29)$$

which results if we set all the numbers $c_{r+1}, \dots, c_n$ in formulas (3.20) equal to zero.

Adding that particular solution to the general solution (3.26). of the corresponding

homogeneous system, we get the following expression for the general solution of the nonhomogeneous system (3.1):

$$X = X_0 + C_1 X_1 + C_2 X_2 + \dots + C_{n-r} X_{n-r}, \quad (3.30)$$

In this expression X_0 is a particular solution (3.29), $C_1, C_2, \dots, C_{n-r}$ are arbitrary constants, and $X_1, X_2, \dots, X_{n-r}$ are elements of the normal fundamental collection of solutions (3.25) of the corresponding homogeneous system.

Thus, for nonhomogeneous system (3.21) considered at the end of 3.2.2, a particular solution of form (3.29) is $X_0 = (6, 2, 0, 0)$. Adding this particular solution to the general solution (3.28) of the corresponding homogeneous system (3.27), we get the following general solution of nonhomogeneous system (3.21):

$$X_0 = (6, 2, 0, 0) + C_1 \left(-\frac{3}{2}, -\frac{1}{2}, 1, 0 \right) + C_2 (-1, -2, 0, 1).$$

(Here C_1 and C_2 are arbitrary constants.) 1

3.2.4. Concluding remarks on solution of linear systems. The methods of solving linear systems discussed in detail in the previous subsections are held up by the necessity of calculating the rank of the matrix and finding its base minor. When the base minor is known, the solution reduces to the technique of calculating the determinants and using Cramer's rule.

The following rule can be used to calculate the rank of a matrix: *in calculating the rank of a matrix, we must pass from minors of lower orders to those of higher orders; and, if we have found the nonzero minor M of order k , then we have only to calculate the minors of order $(k+1)$, bordering* that minor M ; if all the **bordering** minors of order $(k+1)$ are equal to zero, then the rank of the matrix is equal to k .*

Here is another rule for calculating the rank of a matrix. Note that we can do **three simple operations** on the rows (columns) of a matrix without changing the rank of the matrix: (1) transposition of two rows (or two columns), (2) multiplication of a row (or a column) by any nonzero factor, (3) addition to a row (column) of an arbitrary linear combination of the other rows (columns).

We say that the matrix $\|a_{ij}\|$ containing m rows and n columns is of a diagonal form if all the elements different from $a_{11}, a_{22}, \dots, a_{rr}$, where $r = \min \{m, n\}$, are equal to zero. The rank of such a matrix is evidently equal to r .

Let us verify that by means of three simple operations we can reduce any matrix to a diagonal form (and this enables us to calculate its rank).

$$A = \begin{vmatrix} a_{11} & \dots & a_{1n} \\ \dots & \dots & \dots \\ a_{m1} & \dots & a_{mn} \end{vmatrix} \quad (3.31)$$

Indeed, if all the elements of matrix (3.31) are equal to zero, this means that the matrix has been reduced to a diagonal form. Now if matrix (3.31) has nonzero elements, then by interchanging two rows and two columns we can make the element a_{11} differ from zero. Then, multiplying the first row of the matrix by a_{11}^{-1} , we turn the element a_{11} into unity. Subtracting then the first column multiplied by a_{1j} from the j th column of the matrix (for $j = 2, 3, \dots, n$) and after that subtracting the first row multiplied by a_{i1} from the i th row (for $i = 2, 3, \dots, m$), we get a matrix of the form

$$\begin{vmatrix} 1 & 0 & \dots & 0 \\ 0 & a'_{22} & \dots & a'_{2n} \\ \dots & \dots & \dots & \dots \\ 0 & a'_{m2} & \dots & a'_{mn} \end{vmatrix}$$

instead of (3.31). Performing the indicated operations on the matrix in the frame and proceeding as before, we get, after a finite number of steps, a diagonal matrix.

The methods of solving linear systems, using, in the final analysis, the apparatus of Cramer's rule, which we have discussed in the previous subsections, can lead to major errors when the values of the coefficients of the equations and constant terms are approximate or when these values are rounded off in the process of calculation. First and foremost, it concerns the case when a matrix corresponding to the principal determinant (or to the base minor) is **ill-conditioned** (i.e. when "small" changes in the elements of that matrix are associated with "great" changes in the elements of

the inverse matrix). Naturally, the solution of a linear system will be unstable in this case (i.e. “small” changes in the values of the coefficients of the equations and constant terms will be associated with “great” changes in the solutions). These circumstances necessitate an elaboration of other (distinct from Cramer’s rule) theoretical algorithms for seeking solutions and numerical methods of solving linear systems.

In 4.4 we shall get acquainted with Tikhonov’s **regularization method** for seeking the so-called **normal** (i.e. the closest to the origin) solution of a linear system.

Chapter 6 contains the fundamental knowledge of the so-called **iteration methods** of solution of linear systems making it possible to solve such systems by means of successive approximations of the unknowns.

Chapter – 4

EUCLIDEAN SPACES

The reader is sure to be acquainted, from the course of analytic geometry, with the concept of a scalar product of two free vectors, and with the four principal properties of that scalar product. In this chapter we shall study linear spaces of any nature for whose elements in any way (no matter which) a rule is defined associating ‘any two elements with a number called the scalar product of those elements. It is only essential that the rule should possess the same four properties as the rule of forming a scalar product of two free vectors. Linear spaces in which the indicated rule is defined are called **Euclidean spaces**. In this chapter, we elucidate the main properties of arbitrary Euclidean spaces.

4.1. A Real Euclidean Space and Its Simplest Properties

4.1.1. Definition of a real Euclidean space. *A real linear space R is called a **real Euclidean space** (or, simply, a **Euclidean space**) if the following two requirements are met:*

*I. There is a rule which puts any two elements x and y of that space into correspondence with a real number called the **scalar product** of those elements and symbolized as (x, y) .*

II. This rule is subordinated to the following four axioms:

1°. $(x, y) = (y, x)$ (commutative property or symmetry). _

2°. $(x_1 + x_2, y) = (x_1, y) + (x_2, y)$ (distributive property). _

3°. $(\lambda x, y) = \lambda (x, y)$ for any real λ .

4°. $(x, x) > 0$ if x is a nonzero element, $(x, x) = 0$ if x is a zero element.

It should be emphasized that when introducing the notion of a Euclidean space we abstract ourselves not only from the nature of the objects in question“ but also from the specific features of the rules for the formation of the sum of elements, the product of an element by a scalar and the scalar product of elements (it is only

essential that the rules satisfy the eight axioms of a linear space and the four axioms of a scalar product).

Now if the nature of the objects in question and the form of the indicated rules are specified, then the Euclidean space is said to be specific.

Here are some examples of specific Euclidean spaces.

Example 1. Let us consider the linear space B_3 of all free vectors. We define the scalar product of any two vectors as we did it in analytic geometry (i.e. as the product of the lengths of these vectors by the cosine of the angle between them). In the course of analytic geometry, we proved the validity of axioms 1°-4° for a scalar product thus defined. Hence, the space B_3 with a scalar product thus defined is a Euclidean space.

Example 2. Let us consider the finite-dimensional linear space $C[a, b]$ of all functions $x(t)$ defined and continuous on the interval $a \leq t \leq b$. The scalar product of two such functions $x(t)$ and $y(t)$ is defined as an integral (from a to b) of the product of these functions

$$\int_a^b x(t) y(t) dt. \quad (4.1)$$

The validity of axioms 1° - 4° for the scalar product thus defined can be simply verified. Indeed, the validity of axiom 1° is obvious, the validity of axioms 2° and 3° follows from the linear properties of a definite integral; the validity of axiom 4°

follows from the fact that the integral $\int_a^b x^2(t) dt$ of the continuous nonnegative

function $x^2(t)$ is nonnegative and vanishes only if that function is, identically zero on the interval $a \leq t \leq b$ (i.e. is a zero element of the space being considered). Thus, the space $C[a, b]$ with the scalar product thus defined is an infinite-dimensional Euclidean space.

Example 3. This example of a Euclidean space gives an n -dimensional linear space A^n of ordered collections of n real numbers, the scalar product of whose any two elements $\mathbf{x} = (x_1, x_2, \dots, x_n)$ and $\mathbf{y} = (y_1, y_2, \dots, y_n)$ is defined by the equation

$$(\mathbf{x}, \mathbf{y}) = x_1 y_1 + x_2 y_2 + \dots + x_n y_n. \quad (4.2)$$

The validity of axiom 1° for the scalar product thus defined is obvious, the validity of axioms 2° and 3° can be easily verified (it is sufficient to recall the definition of the operations of addition of elements and their multiplication by scalars:

$$(x_1 x_2, \dots, x_n) + (y_1, y_2, \dots, y_n) = (x_1 + y_1, x_2 + y_2, \dots, x_n + y_n),$$

$$\lambda(x_1 x_2, \dots, x_n) = (\lambda x_1, \lambda x_2, \dots, \lambda x_n);$$

and, finally, the validity of axiom 4° follows from the fact that $(\mathbf{x}, \mathbf{x}) = x_1^2 + x_2^2 + \dots + x_n^2$ is always a nonnegative number and only under the condition $x_1 = x_2 = \dots = x_n = 0$.

The Euclidean space considered in this example is often symbolized as E^n .

Example 4. In the same linear space A^n we introduce a scalar product of any two elements $\mathbf{x} = (x_1, x_2, \dots, x_n)$ and $\mathbf{y} = (y_1, y_2, \dots, y_n)$ not by relation (4.2) but in some other, more general, way.

For that purpose, we consider a square matrix of order n :

$$A = \begin{vmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \cdot & \cdot & \cdot & \cdot \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{vmatrix} \quad (4.3)$$

Let us form, with the aid of matrix (4.3), a homogeneous polynomial of order 2 with respect to n variables $x_1, x_2, \dots, x_n$

$$\sum_{i=1}^n \sum_{k=1}^n a_{ik} x_i x_k. \quad (4.4)$$

Foretelling events, we note that such a polynomial is called a **quadratic form** (generated by matrix (4.3)).

Quadratic form (4.4) is said to be **positive definite** if it assumes strictly positive values for all the values of the variables $x_1, x_2, \dots, x_n$ which are simultaneously nonzero**. Since for $x_1 = x_2 = \dots = x_n = 0$ the quadratic form (4.4) is evidently equal to zero, we can say that a *positive definite quadratic form vanishes only under the condition* $x_1 = x_2 = \dots = x_n = 0$.

Let us require that matrix (4.3) satisfy the following two conditions:

- 1°. It should generate a positive definite quadratic form (4.4).
- 2°. It should be symmetric (about the principal diagonal), i.e. satisfy the condition $a_{ik} = a_{ki}$ for all $i = 1, 2, \dots, n$ and all $k = 1, 2, \dots, n$.

Let us use matrix (4.3) satisfying conditions 1° and 2° to define the scalar product of any two elements $\mathbf{x} = (x_1, x_2, \dots, x_n)$ and $\mathbf{y} = (y_1, y_2, \dots, y_n)$ of the space A^n by the relation

$$(\mathbf{x}, \mathbf{y}) = \sum_{i=1}^n \sum_{k=1}^n a_{ik} x_i y_k. \quad (4.5)$$

It is easy to verify the validity of axioms 1°-4° for the scalar product thus defined. Indeed, Axioms 2° and 3° are, evidently, valid for arbitrary matrix (4.3), the validity of axiom 1° follows from the condition of symmetry of matrix (4.3), and the validity of axiom 4° follows from the fact that the quadratic form (4.4), which is a scalar product $(\mathbf{x}, \mathbf{x})$, is positive definite.

Thus, the space A^n with a scalar product defined by equation (4.5) is a Euclidean space under the condition of symmetry of matrix (4.3) and positive definiteness of the quadratic form it generates.

If we take an identity matrix as matrix (4.3), then relation (4.4) turns into (4.2) and we get the Euclidean space E^n considered in Example 3.

4.1.2. The simplest properties of an arbitrary Euclidean space. The properties that will be established here are valid for an arbitrary Euclidean space of both finite and infinite dimensionality.

Theorem 4.1. *The inequality*

$$(\mathbf{x}, \mathbf{y})^2 \leq (\mathbf{x}, \mathbf{x})(\mathbf{y}, \mathbf{y}) \quad (4.6)$$

known as the **Cauchy-Bunyakovsky inequality**, is valid for any two elements $\mathbf{x}$ and $\mathbf{y}$ of an arbitrary Euclidean space.

Proof. The inequality $(\lambda\mathbf{x} - \mathbf{y}, \lambda\mathbf{x} - \mathbf{y}) \geq 0$ is valid for any real number λ by virtue of axiom 4° of a scalar product. According to axioms 1° - 3°, we can rewrite the last inequality as

$$\lambda^2 (\mathbf{x}, \mathbf{x}) - 2\lambda (\mathbf{x}, \mathbf{y}) + (\mathbf{y}, \mathbf{y}) \geq 0.$$

The necessary and sufficient condition for the last quadratic trinomial to be nonnegative is the non-positiveness of its discriminant, i.e. the inequality

$$(\mathbf{x}, \mathbf{y})^2 - (\mathbf{x}, \mathbf{x})(\mathbf{y}, \mathbf{y}) \leq 0.$$

Inequality (4.6) follows immediately from (4.7). We have proved the theorem.

Our immediate task is to introduce the notion of the **norm** (or **length**) of every element in an arbitrary Euclidean space. For that purpose, we introduce the notion of a normed linear space.

Definition. The linear space R is said to be **normed** if the following two requirements are satisfied:

I. There is a rule which puts every element $\mathbf{x}$ of the space R into correspondence with a real number called the **norm** (or the **length**) of the indicated element and is designated as $\|\mathbf{x}\|$.

II. The indicated rule is subordinated to the following three axioms:

1°. $\|\mathbf{x}\| > 0$ if $\mathbf{x}$ is a nonzero element; $\|\mathbf{x}\| = 0$ if $\mathbf{x}$ is a zero element.

2°. $\|\lambda\mathbf{x}\| = |\lambda| \|\mathbf{x}\|$ for any element $\mathbf{x}$ and any real number λ .

3°. The following inequality is valid for any two elements $\mathbf{x}$ and $\mathbf{y}$:

$$\|\mathbf{x} + \mathbf{y}\| \leq \|\mathbf{x}\| + \|\mathbf{y}\|, \quad (4.8)$$

called **triangle inequality** (or **Minkowski's inequality**).

Theorem 4.2. Every Euclidean space is normed if the norm of its any element $\mathbf{x}$ is defined by the equation

$$\|x\| = \sqrt{(x, x)}. \quad (4.9)$$

Proof. It is sufficient to prove that axioms 1°-3° from the definition of a normed space hold true for the norm defined by relation (4.9).

The validity of axiom 1° for the norm follows immediately from axiom 4° for a scalar product. The validity of axiom 2° for the norm follows almost directly from axioms 1° and 3° for a scalar product.

It remains to verify the validity of axiom 3°, i.e. inequality (4.8), for the norm. We proceed from the Cauchy-Bunyakovsky inequality (4.6), which we rewrite as

$$|(x, y)| \leq \sqrt{(x, x)} \sqrt{(y, y)}. \quad (4.7')$$

With the aid of the last inequality, axioms 1° - 4° for a scalar product and the definition of a norm, we obtain

$$\begin{aligned} \|x + y\| &= \sqrt{(x + y, x + y)} = \sqrt{(x, x) + 2(x, y) + (y, y)} \\ &\leq \sqrt{(x, x) + \sqrt{(x, x)} \cdot \sqrt{(y, y)} + (y, y)} \\ &= \sqrt{[\sqrt{(x, x)} + \sqrt{(y, y)}]^2} \end{aligned}$$

And this completes the proof of the theorem.

Corollary. *In every Euclidean space with the norm of the elements defined by relation (4.9), the triangle inequality (4.8) is valid for any two elements x and y .*

Note that in any real Euclidean space the notion of the angle between two arbitrary elements x and y of the space can be introduced. By complete analogy with vector algebra, we take the angle (varying from 0 to π) whose cosine is defined by the relation

$$\cos \Phi = \frac{(x, y)}{\|x\| \|y\|} = \frac{(x, y)}{\sqrt{(x, x)} \sqrt{(y, y)}}$$

as the angle Φ between the elements x and y .

The definition of the angle we have given is correct since, by virtue of the Cauchy-Bunyakovsky inequality (4.7'), the fraction on the right-hand side of the last equation does not exceed unity in absolute value.

Next we agree to call two arbitrary elements $\mathbf{x}$ and $\mathbf{y}$ of the Euclidean space E *orthogonal* if the scalar product of those elements $(\mathbf{x}, \mathbf{y})$ is zero (in that case the angle ϕ between the elements $\mathbf{x}$ and $\mathbf{y}$ is zero).

Appealing again to vector algebra, we call the sum $\mathbf{x} + \mathbf{y}$ of two **orthogonal elements** $\mathbf{x}$ and $\mathbf{y}$ the **hypotenuse** of a right triangle constructed on the elements $\mathbf{x}$ and $\mathbf{y}$.

Note that the **Pythagoras theorem** holds true for every Euclidean space: the square of the hypotenuse is equal to the sum of the squares of the lengths of the legs. Indeed, since $\mathbf{x}$ and $\mathbf{y}$ are orthogonal and $(\mathbf{x}, \mathbf{y}) = 0$, we have

$$\begin{aligned}\|\mathbf{x} + \mathbf{y}\|^2 &= (\mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y}) = (\mathbf{x}, \mathbf{x}) + 2(\mathbf{x}, \mathbf{y}) + (\mathbf{y}, \mathbf{y}) \\ &= (\mathbf{x}, \mathbf{x}) + (\mathbf{y}, \mathbf{y}) = \|\mathbf{x}\|^2 + \|\mathbf{y}\|^2\end{aligned}$$

by virtue of the axioms and the definition of a norm.

This result can be generalized to n pairwise orthogonal elements

$$\begin{aligned}\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n: \text{ if } \mathbf{z} = \mathbf{x}_1 + \mathbf{x}_2 + \dots + \mathbf{x}_n, \text{ then} \\ \|\mathbf{z}\|^2 &= (\mathbf{x}_1 + \mathbf{x}_2 + \dots + \mathbf{x}_n, \mathbf{x}_1 + \mathbf{x}_2 + \dots + \mathbf{x}_n) \\ &= (\mathbf{x}_1, \mathbf{x}_1) + (\mathbf{x}_2, \mathbf{x}_2) + \dots + (\mathbf{x}_n, \mathbf{x}_n) \\ &= \|\mathbf{x}_1\|^2 + \|\mathbf{x}_2\|^2 + \dots + \|\mathbf{x}_n\|^2.\end{aligned}$$

In conclusion we shall write the norm, the Cauchy-Bunyakovsky inequality and the triangle inequality in each specific Euclidean space discussed in 4.1.1.

In a Euclidean space of all free vectors with the ordinary definition of the scalar product, the norm of the vector $\mathbf{a}$ coincides with its length $|\mathbf{a}|$, the Cauchy-Bunyakovsky inequality can be reduced to the form $(\mathbf{a}, \mathbf{b})^2 \leq |\mathbf{a}|^2 |\mathbf{b}|^{2*}$ and the triangle inequality can be reduced to the form $|\mathbf{a} + \mathbf{b}| \leq |\mathbf{a}| + |\mathbf{b}|^{**}$.

*For the scalar product of the vectors $(\mathbf{a}, \mathbf{b}) = |\mathbf{a}| |\mathbf{b}| \cos \phi$ this inequality trivially follows from the fact that $\cos^2 \phi \leq 1$.

** If we add the vectors $\mathbf{a}$ and $\mathbf{b}$ together by the rule of triangle, then this inequality can be trivially reduced to the statement that one side of a triangle does not exceed the sum of its other two sides.

In the Euclidean space $C[a, b]$ of all functions $x = x(t)$ with the scalar product (4.1) continuous on the interval $a \leq t \leq b$, the norm of the element

$x = x(t)$ is $\sqrt{\int_a^b x^2(t) dt}$, and the Cauchy-Bunyakovsky inequality and the triangle

inequality have the forms

$$\left[\int_a^b x(t) y(t) dt \right]^2 \leq \int_a^b x^2(t) dt \int_a^b y^2(t) dt,$$

$$\sqrt{\int_a^b [x(t) y(t)]^2 dt} \leq \sqrt{\int_a^b x^2(t) dt} + \sqrt{\int_a^b y^2(t) dt}.$$

These two inequalities play a significant role in various divisions of mathematical analysis.

In the Euclidean space E^n of ordered collections of n real numbers with scalar product (4.2) the norm of any element $x = (x_1, x_2, \dots, x_n)$

$$\|x\| = \sqrt{x_1^2 + x_2^2 + \dots + x_n^2}.$$

and the Cauchy-Bunyakovsky inequality and the triangle inequality have the forms

$$(x_1 y_1 + x_2 y_2 + \dots + x_n y_n)^2$$

$$\leq (x_1^2 + x_2^2 + \dots + x_n^2) (y_1^2 + y_2^2 + \dots + y_n^2)$$

$$\sqrt{(x_1 + y_1)^2 + (x_2 + y_2)^2 + \dots + (x_n + y_n)^2}$$

$$\leq \sqrt{x_1^2 + x_2^2 + \dots + x_n^2} + \sqrt{y_1^2 + y_2^2 + \dots + y_n^2}.$$

Finally, in a Euclidean space of ordered collections of n real numbers with scalar product (4.5), the norm of any element $\mathbf{x} = (x_1, x_2, \dots, x_n)$ is* (* Recall that in this case matrix (4.3) is symmetric and generates a positive definite quadratic form (4.4))

$$\|\mathbf{x}\| = \sqrt{\sum_{i=1}^n \sum_{k=1}^n a_{ik} x_i x_k}$$

and the Cauchy-Bunyakovsky and triangle inequalities have the forms

$$\begin{aligned} \left(\sum_{i=1}^n \sum_{k=1}^n a_{ik} x_i y_k \right)^2 &\leq \left(\sum_{i=1}^n \sum_{k=1}^n a_{ik} x_i x_k \right) \left(\sum_{i=1}^n \sum_{k=1}^n a_{ik} y_i y_k \right) \\ \sqrt{\sum_{i=1}^n \sum_{k=1}^n a_{ik} (x_i + y_i)(x_k + y_k)} &\leq \sqrt{\sum_{i=1}^n \sum_{k=1}^n a_{ik} x_i x_k} \\ &+ \sqrt{\sum_{i=1}^n \sum_{k=1}^n a_{ik} y_i y_k}. \end{aligned}$$

4.2. An Orthonormal Basis of a Finite-Dimensional Euclidean Space

We shall study here Euclidean spaces of a finite dimension n . The extension of the results discussed here to infinite-dimensional Euclidean spaces does not enter the scope of this book and is the subject for special discussion. (Such spaces are discussed in Chapters 10 and 11 of our *Fundamentals of Mathematical Analysis*, Part 2.)

4.2.1. The concept of an orthonormal basis and its existence. In Chapter 2 we introduced the notion of a basis of an n -dimensional linear space. In a linear space all bases are equivalent and we had no reason there to prefer one basis to another.

There are special, particularly convenient, bases in a Euclidean space called **orthonormal** bases. They play the same role as the Cartesian rectangular basis in analytic geometry. Let us define an orthonormal basis.

Definition. We say that n elements $e_1, e_2, \dots, e_n$ of the n -dimensional Euclidean space E form an **orthonormal basis** of that space if these elements are pairwise orthogonal and the norm of each of these elements is equal to unity, i.e. if

$$(e_i \ e_k) = \begin{cases} 1 & \text{if } i = k, \\ 0 & \text{if } i \neq k. \end{cases} \quad (4.10)$$

To see whether this definition is correct, we must prove that the elements $e_1, e_2, \dots, e_n$ entering into this definition form one of the bases of the n -dimensional space E being considered, and for that purpose, by virtue of Theorem 2.5, it is sufficient to prove that the elements $e_1, e_2, \dots, e_n$ are linearly independent, i.e. that the equality

$$\alpha_1 e_1 + \alpha_2 e_2 + \dots + \alpha_n e_n = 0 \quad (4.11)$$

is possible only if $\alpha_1 = \alpha_2 = \dots = \alpha_n = 0$.

Let us prove that this equality holds true. Suppose k is **any** one of the numbers 1, 2, ..., n . Multiplying scalarly equality (4.11) by the element e_k and using the axioms of a scalar product and relations (4.10), we find that $\alpha_k = 0$.

Let us now prove the following *fundamental* theorem.

Theorem 4.3. *There is an orthonormal basis in every n -dimensional Euclidean space E .*

Proof. In accordance with the definition of dimensionality there are n linearly independent elements $f_1, f_2, \dots, f_n$ in the space E .

We shall prove that it is possible to construct n elements $e_1, e_2, \dots, e_n$ which can be linearly expressed in terms of $f_1, f_2, \dots, f_n$ and which form an orthonormal basis (i.e. satisfy relations (4.10)).

Let us prove the possibility of constructing such elements $e_1, e_2, \dots, e_n$ by induction.

If there is only one element f_1 , then, to construct the element e_1 with a norm equal to unity, it is sufficient to normalize the element f_1 , i.e. to multiply that element by the number $[\sqrt{f_1, f_1}]^{-1}$ inverse to its norm*. (* Recall that there cannot be a zero element among the linearly independent elements $f_1, f_2, \dots, f_n$ and so the norm of f_1 exceeds zero.)

As a result, we obtain an element $e_1 = [\sqrt{f_1, f_1}]^{-1} f_1$ with a norm equal to unity.

Supposing that m is an integer smaller than n , we assume that we have constructed m elements $e_1, e_2, \dots, e_m$ which can be linearly expressed in terms of $f_1, f_2, \dots, f_m$ are pairwise orthogonal and have norms equal to unity. Let us prove that we can add to these elements $e_1, e_2, \dots, e_m$ one more element e_{m+1} which can be linearly expressed in terms of $f_1, f_2, \dots, f_{m+1}$, is orthogonal to each of the elements $e_1, e_2, \dots, e_m$ and has a norm equal to unity.

Let us verify that the element e_{m+1} has the form

$$e_{m+1} = \alpha_{m+1} [f_{m+1} - (f_{m+1}, e_m) e_m - (f_{m+1}, e_{m-1}) e_{m-1} - \dots - (f_{m+1}, e_1) e_1], \quad (4.12)$$

where α_{m+1} is some real number.

Indeed, the element e_{m+1} can be linearly expressed in terms of $f_1, f_2, \dots, f_{m+1}$ (because it can be linearly expressed in terms of $e_1, e_2, \dots, e_m, f_{m+1}$ and each of the elements $e_1, e_2, \dots, e_m$ can be linearly expressed in terms of $f_1, f_2, \dots, f_m$).

It follows immediately that for $\alpha_{m+1} \neq 0$ the element e_{m+1} is obviously **nonzero** (since otherwise a certain linear combination of the linearly independent elements $f_1, f_2, \dots, f_{m+1}$ would be a zero element in which by virtue of (4.12) the coefficient in f_{m+1} is different from zero).

Next, from the fact that the elements $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_m$ are pairwise orthogonal and have norms equal to unity and from relation (4.12) it immediately follows that the scalar product $(\mathbf{e}_{m+1}, \mathbf{e}_k)$ is equal to zero for any number k equal to $1, 2, m$.

To complete induction, it remains to prove that the number α_{m+1} may be chosen such that the norm of element (4.12) is equal to unity. We have established above that for $\alpha_{m+1} \neq 0$ the element e_{m+1} and, hence, the element appearing in brackets in (4.12) is nonzero.

Therefore, to normalize the element in brackets, we must take the number α_{m+1} inverse to the positive norm of that element in brackets. In that case, the norm of e_{m+1} is unity. We have thus proved the theorem.

The theorem proved leads to the algorithm, carried out step-by-step, for constructing, from the given system of n linearly independent elements $f_1, f_2, \dots, f_n$ a system of n pairwise orthogonal elements $e_1, e_2, \dots, e_n$ the norm of each of which is equal to unity:

$$\begin{aligned} e_1 &= \frac{f_1}{\sqrt{f_1, f_1}}; \\ e_2 &= \frac{g_2}{\sqrt{g_2, g_2}}, \quad \text{where } g_2 = f_2 - (f_2, e_1) e_1; \\ e_3 &= \frac{g_3}{\sqrt{g_3, g_3}}, \quad \text{where } g_3 = f_3 - (f_3, e_2) e_2 - (f_3, e_1) e_1; \\ &\dots\dots\dots \\ e_4 &= \frac{g_n}{\sqrt{g_n, g_n}}, \quad \text{where } g_n = f_n - (f_n, e_{n-1}) e_{n-1} - \dots - (f_n, e_1) e_1. \end{aligned}$$

This algorithm is known as the **procedure of orthogonalization** of the linearly independent elements $f_1, f_2, \dots, f_n$.

Remark. There are, of course, - many orthonormal bases in every n -dimensional Euclidean space E . Indeed, if, for instance, we use the orthogonalization procedure for constructing an orthonormal basis of the same linear independent elements

$f_1, f_2, \dots, f_n$ then, beginning the orthogonalization process with different elements f_k , we come to different orthonormal bases. In 7.7.2 we shall study the connection between various orthonormal bases of the given Euclidean space E .

We can consider as an example of an orthonormal basis the rectangular Cartesian basis of the Euclidean space of all free, vectors or a collection of n elements

$$\begin{aligned} \mathbf{e}_1 &= (1, 0, 0, \dots, 0), \\ \mathbf{e}_2 &= (0, 1, 0, \dots, 0), \\ &\vdots \\ \mathbf{e}_n &= (0, 0, 0, \dots, 1). \end{aligned}$$

of the Euclidean space E^n of all ordered collections of n real numbers with scalar product (4.2).

4.2.2. The properties of an orthonormal basis. Suppose $e_1, e_2, \dots, e_n$ is an arbitrary orthonormal basis of the n -dimensional Euclidean space E , and x and y are two arbitrary elements of that space. Let us find the expression for the scalar product (x, y) of these elements in terms of their coordinates relative to the basis $e_1, e_2, \dots, e_n$.

We designate the coordinates of the elements $\mathbf{x}$ and $\mathbf{y}$ relative to the basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ as $x_1, x_2, \dots, x_n$ and $y_1, y_2, \dots, y_n$ respectively i.e. assume that $\mathbf{x} = x_1\mathbf{e}_1 + x_2\mathbf{e}_2 + \dots + x_n\mathbf{e}_n$, $\mathbf{y} = y_1\mathbf{e}_1 + y_2\mathbf{e}_2 + \dots + y_n\mathbf{e}_n$. Then

$$(\mathbf{x}, \mathbf{y}) = (x_1\mathbf{e}_1 + x_2\mathbf{e}_2 + \dots + x_n\mathbf{e}_n, y_1\mathbf{e}_1 + y_2\mathbf{e}_2 + \dots + y_n\mathbf{e}_n).$$

By virtue of the axioms of a scalar product and relations (4.10) it follows from the last equation that

$$(\mathbf{x}, \mathbf{y}) = \left(\sum_{i=1}^n x_i \mathbf{e}_i, \sum_{k=1}^n y_k \mathbf{e}_k \right) = \sum_{i=1}^n \sum_{k=1}^n x_i y_k (\mathbf{e}_i \mathbf{e}_k)$$

$$= x_1 y_1 + x_2 y_2 + \dots + x_n y_n.$$

Thus, in the final analysis we have

$$(\mathbf{x}, \mathbf{y}) = x_1 y_1 + x_2 y_2 + \dots + x_n y_n. \quad (4.13)$$

Thus, in an orthonormal basis the scalar product of any two elements is equal to the sum of the products of the respective coordinates of those elements.

Let us consider now an arbitrary (non orthonormal) basis $f_1, f_2, \dots, f_n$ in the n -dimensional Euclidean space E and find the expression of the scalar product of two arbitrary elements $\mathbf{x}$ and $\mathbf{y}$ in terms of their coordinates with respect to the indicated basis.

We designate the coordinates of the elements $\mathbf{x}$ and $\mathbf{y}$ with respect to the basis $f_1, f_2, \dots, f_n$ as $x_1, x_2, \dots, x_n$ and $y_1, y_2, \dots, y_n$ respectively, i.e. assume that

$$\mathbf{x} = x_1 f_1 + x_2 f_2 + \dots + x_n f_n, \mathbf{y} = y_1 f_1 + y_2 f_2 + \dots + y_n f_n.$$

Applying the axioms of a scalar product, we get

$$\begin{aligned} (\mathbf{x}, \mathbf{y}) &= (x_1 f_1 + x_2 f_2 + \dots + x_n f_n, \\ & y_1 f_1 + y_2 f_2 + \dots + y_n f_n) = \left(\sum_{i=1}^n x_i f_i, \sum_{k=1}^n y_k f_k \right) \\ &= \sum_{i=1}^n \sum_{k=1}^n x_i y_k (f_i f_k) \end{aligned}$$

Thus, in an arbitrary basis $f_1, f_2, \dots, f_n$ the scalar product of any two elements $\mathbf{x} = x_1 f_1 + x_2 f_2 + \dots + x_n f_n$ and $\mathbf{y} = y_1 f_1 + y_2 f_2 + \dots + y_n f_n$ is defined by the equation

$$(\mathbf{x}, \mathbf{y}) = \sum_{i=1}^n \sum_{k=1}^n a_{ik} x_i y_k,$$

in which the matrix $\|a_{ik}\| (i=1, 2, \dots, n; k=1, 2, \dots, n)$ has the elements $a_{ik} = (f_i, f_k)$.

The last assertion leads to the following result : for the scalar product of any two elements in the given basis $f_1, f_2, \dots, f_n$ of the Euclidean space E to be equal to the

sum of products of the respective coordinates of those elements, it is necessary and sufficient that the basis $f_1, f_2, \dots, f_n$ be orthonormal.

In fact, expression (4.14) passes into (4.13) if and only if the matrix $\|a_{ik}\|$ with the elements $a_{ik} = (f_1, f_2)$ is an identity matrix, i.e. if and only if the relations

$$(f_i, f_k) = \begin{cases} 1 & \text{if } i = k, \\ 0 & \text{if } i \neq k. \end{cases}$$

establishing the orthonormality of the basis $f_1, f_2, \dots, f_n$ are satisfied.

Let us continue with the discussion of the arbitrary orthonormal basis $e_1, e_2, \dots, e_n$ of the n -dimensional Euclidean space E . We shall elucidate the sense of the coordinates of the arbitrary element x with respect to the indicated basis.

We designate the coordinates of the element x with respect to the basis $e_1, e_2, \dots, e_n$ as $x_1, x_2, \dots, x_n$ i.e. assume

$$x = x_1 e_1 + x_2 e_2 + \dots + x_n e_n. \quad (4.15)$$

Next we designate as k any of the numbers $1, 2, \dots, n$ and multiply both sides of (4.5) scalarly by the element e_k . On the basis of the axioms of a scalar product and relations (4.10), we obtain

$$(x, e_k) = \left(\sum_{i=1}^n x_i e_i, e_k \right) = \sum_{i=1}^n x_i (e_i, e_k) = x_k.$$

Thus, the coordinates of an arbitrary element relative to an orthonormal basis are equal to the scalar products of that element by the respective components.

Since the scalar product of the arbitrary element x by the element e with a norm equal to unity is naturally called the **projection of the element x by the element e** , we may say that *the coordinates of an arbitrary element relative to an orthonormal basis are equal to the projections of that element by the respective components.*

Thus, an arbitrary orthonormal basis possesses properties which are quite analogous to the properties of a rectangular Cartesian basis.

4.2.3. Decomposition of an n -dimensional Euclidean space in a direct sum of its subspace and its orthogonal complement. Suppose G is an arbitrary subspace of an n -dimensional Euclidean space E .

*The collection F of all elements y of the space E which are orthogonal to every element x of the subspace G is known as an **orthogonal complement** of the subspace G .*

Note that the orthogonal complement F is itself a subspace of E (since it evidently follows from the orthogonality of every element y_1 and y_2 to the element x that any linear combination of the elements y_1 and y_2 is orthogonal to the element x).

Let us prove that *every n -dimensional Euclidean space E is a direct sum of its arbitrary subspace G and its orthogonal complement F .*

We choose in G an arbitrary orthonormal basis $e_1, e_2, \dots, e_k$. By virtue of what was proved in 2.3.1, this basis can be complemented by the elements $f_{k+1}, \dots, f_n$ of the space E to complete the basis in the whole E . Carrying out the **process of orthogonalization** of elements $f_1, \dots, e_k, f_{k+1}, \dots, f_n$, we get an orthonormal basis $e_1, \dots, e_k, e_{k+1}, \dots, e_n$ of the whole space E . Expressing the arbitrary element x of the space E in terms of these components, i.e. representing it as $x = x_1 e_1 + \dots + x_k e_k + x_{k+1} e_{k+1} + \dots + x_n e_n$, we find that the element x can be uniquely represented in the form $x = x' + x''$, where $x' = x_1 e_1 + \dots + x_k e_k$ is a definite element of G and $x'' = x_{k+1} e_{k+1} + \dots + x_n e_n$ is a definite element of the orthogonal complement F (each element $e_{k+1}, \dots, e_n$ is orthogonal to any of the elements $e_1, \dots, e_k$ and is, therefore, orthogonal to any element of G , the linear combination $x_{k+1} e_{k+1} + \dots + x_n e_n$ is, therefore, also orthogonal to any element of G , i.e. is a definite element of F).

4.2.4. The isomorphism of n -dimensional Euclidean spaces. We shall show here that various Euclidean spaces of the same dimensionality n do not essentially differ from one another in the sense of the properties connected with the operations introduced in those spaces.

Since only operations of addition of elements, multiplication of elements by scalars and a scalar multiplication of elements were introduced in Euclidean spaces, it is natural to formulate the following definition.

Definition. Two Euclidean spaces E and E' are said to be **isomorphic** if a one-to-one correspondence can be established between the elements of these spaces such that if the elements $\mathbf{x}$ and $\mathbf{y}$ of the space E are associated with the elements $\mathbf{x}'$ and $\mathbf{y}'$ of the space E' , then the element $\mathbf{x} + \mathbf{y}$ is associated with the element $\mathbf{x}' + \mathbf{y}'$, the element $\lambda \mathbf{x}$ (for any real λ) is associated with the element $\lambda \mathbf{x}'$ and the scalar product $(\mathbf{x}, \mathbf{y})$ is equal to the scalar product $(\mathbf{x}', \mathbf{y}')$.

Thus, the Euclidean spaces E and E' are isomorphic if they are isomorphic as linear spaces and if that isomorphism retains the value of the scalar product of the corresponding pairs of elements.

Theorem 4.4. All Euclidean spaces of the same dimension n are isomorphic to one another.

Proof. It is sufficient to prove that any n -dimensional Euclidean space E' is isomorphic to the Euclidean space E^n of ordered collections of n real numbers with scalar product (4.2). According to Theorem 4.3, there is an orthonormal basis $\mathbf{e}'_1, \mathbf{e}'_2, \dots, \mathbf{e}'_n$ in the Euclidean space E' . We put every element $\mathbf{x}' = x_1 \mathbf{e}'_1 + x_2 \mathbf{e}'_2 + \dots + x_n \mathbf{e}'_n$ of the space E' into correspondence with n real numbers $x_1, x_2, \dots, x_n$, i.e. a definite element $\mathbf{x} = (x_1, x_2, \dots, x_n)$ of the space E^n .

The correspondence we have established is one-to-one. In addition, it follows from Theorem 2.4 that if the elements $\mathbf{x}' = (x_1, x_2, \dots, x_n)$ and $\mathbf{y}' = (y_1, y_2, \dots, y_n)$ of the space E' are associated with the respective elements $\mathbf{x} = (x_1, x_2, \dots, x_n)$ and $\mathbf{y} = (y_1, y_2, \dots, y_n)$ of the space E^n , then the element $\mathbf{x}' + \mathbf{y}'$ is associated with the element $\mathbf{x} + \mathbf{y}$ and the element $\lambda \mathbf{x}'$ is associated with the element $\lambda \mathbf{x}$.

It remains to prove that the value of the scalar product is retained for the respective pairs of elements $\mathbf{x}', \mathbf{y}'$ and $\mathbf{x}, \mathbf{y}$.

Because of the orthonormality of the basis $e'_1, e'_2, \dots, e'_n$ and according to formula (4.13), we have $(x' + y') = x_1 y_1 + x_2 y_2, \dots, x_n y_n$. On the other hand, by virtue of formula (4.2) defining the scalar product in the space E^n , $(x, y) = x_1 y_1 + x_2 y_2, \dots, x_n y_n$. And this completes the proof of the theorem.

This theorem enables us to assert that if in any specific n -dimensional Euclidean space E' the theorem formulated in terms of operations of addition, multiplication by scalars and scalar multiplication of elements has been proved, then that theorem is valid in an arbitrary n -dimensional Euclidean space E as well.

4.3. A Complex Euclidean Space

4.3.1. Definition of a complex Euclidean space. We indicated at the end of 2.1.1 that if in the definition of a linear space we take the numbers $\lambda, \mu, \dots$ not from the set of real numbers but from the set of all complex numbers, then we arrive at the concept of a complex linear space.

A complex Euclidean space which is of a fundamental importance in the theory of nonself-adjoint linear transformations is constructed on the basis of complex linear spaces.

To introduce a complex linear space we must first introduce the notion of a scalar product of two elements of a complex linear space subordinated to the corresponding four axioms.

Definition. *The complex linear space R is called a **complex Euclidean space** if the following two requirements are met:*

I. *There is a rule which puts any two elements x and y of that space into correspondence with a **complex** number called the **scalar product** of those elements and symbolized as (x, y) .*

II. *The indicated rule is subordinated to the following four axioms*

$$1^\circ. (x, y) = \overline{(y, x)}^*.$$

$$2^\circ. (x_1 + x_2, y) = (x_1, y) + (x_2, y).$$

$$3^\circ. (\lambda x, y) = \lambda (x, y).$$

4°. (x, x) is a real nonnegative number vanishing only when x is a zero element**.

The following two relations are logical consequences of axioms 1° - 3°:

$$(x, \lambda y) = \overline{\lambda} (x, y), (x, y_1 + y_2) = (x, y_1) + (x, y_2).$$

Indeed, we infer from axioms If and 3° that

$$(x, \lambda y) = \overline{(\lambda y, x)} = \overline{\lambda(y, x)} = \overline{\lambda} (x, y),$$

and from axioms 1° and 2° we find that

$$(x, y_1 + y_2) = \overline{(y_1 + y_2, x)} = \overline{(y_1, x)} + \overline{(y_2, x)} = (x, y_1) + (x, y_2).$$

Here are some examples of specific complex Euclidean spaces.

Example 1. Let us consider the collection $C^*[a, b]$ of all functions $z = z(t)$ defined for the values of t from the interval $a \leq t \leq b$ and assuming complex values $z(t) = x(t) + iy(t)$ such that the real functions $x(t)$ and $y(t)$ are continuous on that interval. The operations of addition of those functions and their multiplication by complex numbers will be taken from the analysis. The scalar product of any two functions of this kind is defined by the relation $(z_1(t), z_2(t)) = \int_a^b z_1(t) \overline{z_2(t)} dt$.

It is easy to verify all the axioms 1° - 4° for a scalar product thus defined, from which it follows that the collection being considered is a complex Euclidean space.

Example 2. Let us consider the complex linear space A^n whose elements are ordered collections of n complex numbers $x_1, x_2, \dots, x_n$ with the definitions of operations of addition of elements and their multiplication by scalars as in the case of the real linear space A^n .

The scalar product of any two elements $x = (x_1, x_2, \dots, x_n)$ and

$y = (y_1, y_2, \dots, y_n)$ is defined by the relation

$$(x, y) = x_1 \overline{y_1} + x_2 \overline{y_2} + \dots + x_n \overline{y_n}. \quad (4.16)$$

It is easy to verify the validity of axioms 1° - 3° for the scalar product thus defined. The validity of axiom 4° follows from the relation

$$(\mathbf{x}, \mathbf{x}) = x_1 \bar{x}_1 + x_2 \bar{x}_2 + \dots + x_n \bar{x}_n = |x_1|^2 + |x_2|^2 + \dots + |x_n|^2.$$

This means that the space A with the scalar product (4.16) is a complex Euclidean space.

Example 3. We can introduce the scalar product in the same

complex linear space A_*^n not by relation (4.16) but by a more general relation

$$(\mathbf{x}, \mathbf{y}) = \sum_{i=1}^n \sum_{k=1}^n a_{ik} x_i \bar{y}_k, \quad (4.17)$$

in which $\|a_{ik}\|$ is an arbitrary matrix consisting of complex numbers a_{ik} which satisfy the condition $a_{ik} = \bar{a}_{ki}$, such that the quadratic form $|\sum_{i=1}^n \sum_{k=1}^n a_{ik} x_i \bar{x}_k|$ assumes nonnegative real values for all complex numbers $x_1, x_2, \dots, x_n$, and vanishes only under the condition $|x_1|^2 + |x_2|^2 + \dots + |x_n|^2 = 0$.

We invite the reader to verify that the scalar product thus defined satisfies axioms 1° - 4°.

4.3.2. The Cauchy-Bunyakovsky inequality. The concept of a norm. *Let us prove that the **Cauchy-Bunyakovsky inequality** holds true for any two elements x and y of an arbitrary complex Euclidean space.*

$$|(\mathbf{x}, \mathbf{y})|^2 \leq (\mathbf{x}, \mathbf{x})(\mathbf{y}, \mathbf{y}). \quad (4.18)$$

On the basis of axiom 4° the inequality

$$(\lambda \mathbf{x} - \mathbf{y}, \lambda \mathbf{x} - \mathbf{y}) \geq 0 \quad (4.19)$$

is valid for any complex number λ . Since by virtue of axioms 1° - 3° and their logical consequences we have

$$(\lambda \mathbf{x} - \mathbf{y}, \lambda \mathbf{x} - \mathbf{y}) = \lambda \bar{\lambda} (\mathbf{x}, \mathbf{x}) - \lambda (\mathbf{x}, \mathbf{y}) - \bar{\lambda} (\mathbf{y}, \mathbf{x}) + (\mathbf{y}, \mathbf{y})$$

$$= |\lambda|^2 (x, x) - y, \lambda (x, y) - \bar{\lambda} (\overline{x, y}) + (x, y),$$

inequality (4.19) assumes the form

$$|\lambda|^2 (x, x) - \lambda (x, y) - \bar{\lambda} (\overline{x, y}) + (y, y) \geq 0 \quad (4.20)$$

Let us denote the **argument** of the complex number (x, y) by Φ and represent this number in **trigonometric** form

$$(x, y) = |(x, y)| (\cos \Phi + i \sin \Phi), \quad (4.21)$$

Now we set the complex number

$$\lambda = (\cos \Phi - i \sin \Phi), \quad (4.22)$$

where t is an **arbitrary real number**. It is evident from relations (4.21) and (4.22) that $|\lambda| = |t|$, $\lambda (x, y) = (x, y) = \bar{\lambda} (\overline{x, y}) = t |\overline{(x, y)}|$. Therefore, for λ we have chosen, inequality (4.20) turns into an inequality

$$t^2 (x, x) - 2t |(x, y)| + (y, y) \geq 0. \quad (4.23)$$

which is valid for any real t . The necessary and sufficient condition for the quadratic trinomial on the left-hand side of (4.23) to be nonnegative is the non-positivity of its discriminant, i.e. the inequality $|(x, y)|^2 - (x, y)(y, y) \leq 0$ which is equivalent to inequality (4.18).

We can establish with the aid of the Cauchy-Bunyakovsky inequality (4.18) and arguments completely analogous to the proof of Theorem 4.2 that *every complex Euclidean space is normed if the norm of any element x in it is defined by the relation*

$$\|x\| = \sqrt{(x, x)}. \quad (4.24)$$

*In particular, in any complex Euclidean space with the norm defined by relation (4.24) there holds a **triangle inequality** $\|x + y\| \leq \|x\| + \|y\|$.*

Remark. It should be emphasized that the notion of the angle Φ between two arbitrary elements x and y introduced for a real Euclidean space loses sense for a **complex** Euclidean space (because the scalar product (x, y) is, in general, a complex number).

4.3.3. An orthonormal basis and its properties. The elements $\mathbf{x}$ and $\mathbf{y}$ of an arbitrary complex Euclidean space are said to be **orthogonal** if the scalar product $(\mathbf{x}, \mathbf{y})$ of these elements is equal to zero. The orthonormal basis of an n -dimensional complex Euclidean space is the collection of its elements $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ satisfying the relations

$$(\mathbf{e}_i, \mathbf{e}_k) = \begin{cases} 1 & \text{for } i = k, \\ 0 & \text{for } i \neq k. \end{cases} \quad (4.25)$$

(i.e. pairwise zorthogonal and having norms equal to unity).

We can prove in the same way as in 2.1 that these elements are linearly independent and, therefore, form a basis.

By a complete analogy with the proof of Theorem 4.3 (i.e. with the aid of the orthogonalization procedure) we establish the existence of an orthonormal basis in an arbitrary n -dimensional complex Euclidean space.

Let us express the scalar product of two arbitrary elements $\mathbf{x}$ and $\mathbf{y}$ of an n -dimensional complex Euclidean space in terms of their coordinates $(x_1, x_2, \dots, x_n)$ and $(y_1, y_2, \dots, y_n)$ relative to the orthonormal basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$. Since

$$\mathbf{x} = x_1 \mathbf{e}_1 + x_2 \mathbf{e}_2 + \dots + x_n \mathbf{e}_n, \mathbf{y} = y_1 \mathbf{e}_1 + y_2 \mathbf{e}_2 + \dots + y_n \mathbf{e}_n,$$

we obtain, by virtue of axioms 1° - 4° and relations (4.25),

$$\begin{aligned} (\mathbf{x}, \mathbf{y}) &= \left(\sum_{i=1}^n x_i \mathbf{e}_i, \sum_{k=1}^n y_k \mathbf{e}_k \right) \\ &= \sum_{i=1}^n \sum_{k=1}^n x_i \bar{y}_k (\mathbf{e}_i, \mathbf{e}_k) = x_1 \bar{y}_1 + x_2 \bar{y}_2 + \dots + x_n \bar{y}_n. \end{aligned}$$

Next we express the coordinates $x_1, x_2, \dots, x_n$ of the arbitrary element $\mathbf{x}$ relative to the orthonormal basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$.

Multiplying the expression of this element in terms of the components $\mathbf{x} = x_1 \mathbf{e}_1 + x_2 \mathbf{e}_2 + \dots + x_n \mathbf{e}_n$ scalarly by $\mathbf{e}_k$ and using relations (4.25), we obtain (for any k equal to 1, 2, ..., n)

$$(\mathbf{x}, \mathbf{e}_k) = \left(\sum_{i=1}^n x_i \mathbf{e}_i, \mathbf{e}_k \right) = \sum_{i=1}^n x_i (\mathbf{e}_i, \mathbf{e}_k) = x_k.$$

Thus, as in the case of a real Euclidean space, *the coordinates of the arbitrary element $\mathbf{x}$ with respect to an orthonormal basis are equal to the scalar products of that element by the corresponding base elements.*

By a complete analogy with the proof of Theorem 4.4 we establish that *all complex Euclidean spaces of the same dimension n are isomorphic to each other.*

4.4. The Regularization Method for Seeking a Normal Solution of a Linear system_

We return to the general linear system of m equations in n unknowns of the form (3.1). We write this system in short, in matrix form.

$$AX = B. \quad (4.26)$$

Recall that in this notation the symbol A denotes the matrix $A = \|a_{ij}\|$ ($i = 1, 2, \dots, m; j = 1, 2, \dots, n$), and the symbols X and B denote the columns (or vectors) of the form the first of which must be found and the second is given.

$$X = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}, \quad B = \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_m \end{pmatrix}$$

We shall discuss the case when the values of the elements of the matrix A and of the column of constant terms B are given only approximately. Then it is natural to speak only of approximate values of the sought-for column X . The algorithms for calculating the column of solutions X presented in the previous chapter and based on Cramer's rule may lead to large errors in this case and lose practical significance.

We shall give here an algorithm belonging to Tikhonov and making it possible to find the so-called **normal** (i.e. the closest to the origin) solution of X with an accuracy corresponding to the accuracy of specifying the elements of the matrix A and the column B .

Let us introduce the so-called *spherical norms* of the columns B and X of the matrix A setting them equal to

$$\begin{aligned}\|B\| &= \sqrt{\sum_{i=1}^m b_i^2}, \quad \|X\| = \sqrt{\sum_{j=1}^n x_j^2}, \\ \|A\| &= \sqrt{\sum_{i=1}^m \sum_{j=1}^n a_{ij}^2}.\end{aligned}\tag{4.27}$$

Note that the norms of the columns B and X are defined as ordinary norms of vectors, i.e. the elements of the spaces E^m and E^n respectively. The norm of the matrix A is defined coherently with the norm of the n -dimensional column X in the sense that the norm of the m -dimensional column AX , equal to the product of the matrix A by the column X , satisfies the condition

$$\|AX\| \leq \|A\| \|X\|.\tag{4.28}$$

We assume that instead of exact values of the elements of; the matrix $A = \|a_{ij}\|$ and the column of the right-hand sides $B = \|b_i\|$ we are given approximate values $\tilde{A} = \|\tilde{a}_{ij}\|$, $\tilde{B} = \|\tilde{b}_i\|$.

We call the matrix $\tilde{A}$ (the column $\tilde{B}$) the **δ th approximation** of the matrix A (column B) if the following inequality holds true:

$$\begin{aligned}\|A - \tilde{A}\| &= \sqrt{\sum_{i=1}^m \sum_{j=1}^n (a_{ij} - \tilde{a}_{ij})^2} < \delta \\ \left(\|B - \tilde{B}\| = \sqrt{\sum_{i=1}^m (b_i - \tilde{b}_i)^2} \leq \delta \right).\end{aligned}\tag{4.29}$$

The **normal solution** of the consistent system (4.26) is assumed to be its solution

$$X^0 = \begin{pmatrix} x_1^0 \\ x_2^0 \\ \vdots \\ x_n^0 \end{pmatrix} \quad \text{whose norm } \|X^0\| \text{ is the least of the norms } \|X\| \text{ of all the}$$

solutions X of that system. Note that every consistent system (4.26) (an indeterminate system inclusive) possesses a unique normal solution.

We introduce into consideration the following function of n variables $x_1, x_2, \dots, x_n$

$$\text{or of one column } X = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}$$

$$F^\alpha(x_1, \dots, x_n, \tilde{A}, \tilde{B}) = \|\tilde{A}X - \tilde{B}\|^2 + \alpha \|X\|^2 \quad (4.30)$$

dependent on the elements of the matrix $\tilde{A}$ and the column $\tilde{B}$ as

parameters and also on a certain numerical parameter α . This function can be written out as

$$F^\alpha(X, \tilde{A}, \tilde{B}) = \sum_{i=1}^m \left[\sum_{j=1}^n \tilde{a}_{ij} x_j - \tilde{b}_i \right]^2 + \alpha \sum_{j=1}^n x_j^2. \quad (4.30')$$

Actually $F^\alpha(X, \tilde{A}, \tilde{B})$ is a function of the elements of X of the Euclidean space E^n of the n -dimensional columns. A function of this kind, whose argument are the elements of a certain linear space, is customarily called a **functional**.

It is easy to verify that *for any fixed $\alpha > 0$, the nonnegative functional (4.30')*

$$\text{attains its minimum (in the whole space } E^n \text{) value at the only point } X^\alpha = \begin{pmatrix} x_1^\alpha \\ \vdots \\ x_n^\alpha \end{pmatrix}$$

of the space E^n .

Indeed, differentiating function (4.30') twice, we get

$$\frac{\partial^2 F^\alpha}{\partial x_k \partial x_i} = 2 \sum_{i=1}^m \bar{a}_{ik} a_{il} + 2\alpha \delta_{kl}, \quad \text{where } \delta_{kl} = \begin{cases} 1 & \text{for } k = l, \\ 0 & \text{for } k \neq l. \end{cases}$$

Consequently, the second differential of the function F^α has the form

$$\begin{aligned} d^2 F^\alpha &= \sum_{k=1}^n \sum_{l=1}^n \left[\sum_{i=1}^m \bar{a}_{ik} \bar{a}_{il} \right] dx_k dx_l + \alpha \sum_{k=1}^m \sum_{l=1}^m \delta_{kl} dx_k dx_l \\ &= \sum_{i=1}^m \left[\sum_{k=1}^n \bar{a}_{ik} dx_k \right]^2 + \alpha \sum_{k=1}^n (dx_k)^2. \end{aligned}$$

This equation yields an estimation $d^2 F^\alpha \geq \alpha \sum_{k=1}^n (dx_k)^2$ signifying that the

function F^α is **strictly convex downwards**. In addition,

$F^\alpha \rightarrow \infty$ for $\|X\| \sqrt{\sum_{k=1}^n x_k^2} \rightarrow \infty$. From this it obviously follows that F^α has a

point of minimum X^α and that this point is unique.

The methods of seeking the minimum value of the functionals of the form (4.30) are well elaborated.

Let us prove the following fundamental theorem reducing the problem of an approximate seeking the normal solution of system (4.26) to finding the element

$X^\alpha = \begin{pmatrix} x_1^\alpha \\ \vdots \\ x_n^\alpha \end{pmatrix}$ on which functional (4.30) attains its minimum value.

Theorem of Tikhonov. Suppose the matrix A and the column B satisfy the

conditions ensuring the consistency of system (4.26), $X^0 = \begin{pmatrix} x_1^0 \\ \vdots \\ x_n^0 \end{pmatrix}$ is a normal

solution of that system, $\bar{A}$ is the δ -approximation of the matrix A , B is the δ -approximation of the column B , $\varepsilon(\delta)$ and $\theta(\delta)$ are some increasing functions δ tending to zero as $\delta \rightarrow 0+0$ and such that $\delta^2 \leq \varepsilon(\delta) \alpha(\delta)$. Then, for any $\varepsilon > 0$ there is a positive number $\delta_0 = \delta_0(\varepsilon, \|X^0\|)$ such that for any $\delta < \delta_0(\varepsilon, \|X^0\|)$ and any α satisfying the condition

$$\frac{1}{\varepsilon(\delta)} \delta^2 \leq \alpha \leq \alpha(\delta) \quad (4.31)$$

the element $X^\alpha = \begin{pmatrix} x_1^\alpha \\ \vdots \\ x_n^\alpha \end{pmatrix}$ furnishing functional (4.30) with a minimum satisfies the inequality

$$\|X^\alpha - X^0\| \leq \varepsilon. \quad (4.32)$$

Proof. Let us consider in the linear space E^m a subset $\{U_{\tilde{A}}\}$ of all elements

$$U = \begin{pmatrix} u_1 \\ \vdots \\ u_m \end{pmatrix} \text{ which can be represented as } U = \tilde{A} X, \text{ where } X = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \text{ is an}$$

arbitrary element of the space E^n . It is quite evident that the subset $\{U_{\tilde{A}}\}$ is a linear space and is therefore, a subspace of E^m . We denote by $\{V_{\tilde{A}}\}$ an orthogonal complement of $\{U_{\tilde{A}}\}$ (to complete E^m) and decompose E^m in a direct sum of the subspaces $\{U_{\tilde{A}}\}$ and $\{V_{\tilde{A}}\}$. Assume that $\tilde{B}_{\tilde{A}}$ denotes a (projection of the column $\tilde{B}$ onto the subspace $\{U_{\tilde{A}}\}$ so that $\tilde{B} = \tilde{B}_{\tilde{A}} + (\tilde{B} - \tilde{B}_{\tilde{A}})$, where $(\tilde{B} - \tilde{B}_{\tilde{A}})$ is an element of $\{V_{\tilde{A}}\}$. Then, since for any element X of the space E^n the column $\tilde{A}X$ is an element of $\{U_{\tilde{A}}\}$, we get the following decomposition:

$$\tilde{A}X - \tilde{B} = (\tilde{A}X - \tilde{B}_{\tilde{A}}) + (\tilde{B}_{\tilde{A}} - \tilde{B}),$$

in which the elements $(\tilde{A}X - \tilde{B}_{\tilde{A}})$ and $(\tilde{B}_{\tilde{A}} - \tilde{B})$ are orthogonal to each other and belong to $\{U_{\tilde{A}}\}$ and $\{V_{\tilde{A}}\}$ respectively.

Applying Pythagoras theorem (see 4.1.2), we get (for any element X of the space E^n)

$$\|\tilde{A}X - \tilde{B}\|^2 = \|\tilde{A}X - \tilde{B}_{\tilde{A}}\|^2 + \|\tilde{B}_{\tilde{A}} - \tilde{B}\|^2, \quad (4.33)$$

which yields, in particular, an inequality

$$\|\tilde{B} - \tilde{B}_{\tilde{A}}\| \leq \|\tilde{A}X - \tilde{B}\| \quad (4.34)$$

which is also valid for any element X of the space E^n .

From (4.83) and (4.30) we find that for any X from E^n

$$F^\alpha(X, \tilde{A}, \tilde{B}_{\tilde{A}}) \leq \|\tilde{B} - \tilde{B}_{\tilde{A}}\|^2 F^\alpha(X, \tilde{A}, \tilde{B}_{\tilde{A}}), \quad (4.35)$$

i.e. the functionals on the left-hand and right-hand sides of (4.35) *have a common element X^α furnishing them with a minimum.*

Let us find now for any on satisfying conditions (4.31) the inequality

$$F^\alpha(X^\alpha, \tilde{A}, \tilde{B}_{\tilde{A}}) \leq \alpha \varepsilon |(\delta) C^2 + \alpha \|X^0\| \quad (4.36)$$

in which C denotes the quantity $C = 2(1 + \|X^0\|)$, and X^0 is a normal solution of system (4.26).

Since the column X^α furnishes the functional appearing on the right-hand side of (4.35) with a minimum, we have,

$$F^\alpha(X^\alpha, \tilde{A}, \tilde{B}_{\tilde{A}}) \leq F^\alpha(X^0, \tilde{A}, \tilde{B}_{\tilde{A}}) = \|\tilde{A}X^0 - \tilde{B}_{\tilde{A}}\|^2 + \alpha \|X^0\|^2. \quad (4.37)$$

Using the relation $\tilde{A}X^0 = B$ and the triangle inequality, we obtain

$\|\tilde{A}X^0 - \tilde{B}_{\tilde{A}}\| \leq \|\tilde{A}X^0 - AX^0\| + \|B - \tilde{B}\| + \|\tilde{B} - \tilde{B}_{\tilde{A}}\|$. On the right-hand side of the last inequality, we use relations (4.28) and (4.29) and also inequality (4.34) taken for $X = X^0$. We get

$$\|\tilde{A}X^0 - \tilde{B}_{\tilde{A}}\| \leq \delta \|X^0\| + \delta + \|\tilde{B} - \tilde{A}X^0\|. \quad (4.38)$$

Taking into account once again that $AX^0 = B$ and using the triangle inequality and relations (4.28) and (4.29), we find that

$$\|\tilde{B} - \tilde{A}X^0\| \leq \|\tilde{B} - B\| + \|AX^0 - \tilde{A}X^0\| \leq \delta + \delta \|X^0\| \quad (4.39)$$

It follows from (4.38) and (4.39) that

$$\|\tilde{A}X^0 - \tilde{B}_{\tilde{A}}\| \leq 2\delta (1 + \|X^0\|) = C\delta. \quad (4.40)$$

where $C = 2(1 + \|X^0\|)$.

To complete the proof of estimate (4.36), it remains to substitute (4.40) into (4.37) and use inequality (4.31)

Since it immediately follows from the definition of the functional F^α that $\alpha \|X^\alpha\|^2 \leq F^\alpha(X^\alpha, \tilde{A}, \tilde{B}_{\tilde{A}})$, inequality (4.36) we have proved also yields an inequality

$$\|X^\alpha\| \leq \|X^0\| + \varepsilon_1(\delta), \quad (4.41)$$

in which $\varepsilon_1(\delta) \rightarrow 0$ as $\delta \rightarrow 0+0$. It follows from (4.41) that for all sufficiently small δ the set $\{X^\alpha\}$ of the points X^α of the space E^n is **bounded**.

It is now easy to prove the theorem by *reductio ad absurdum*. We assume that for a certain $\varepsilon_0 > 0$ there is a sequence $\delta_n \rightarrow 0+0$ and a sequence $\{\alpha_n\}$ of numbers corresponding to it which satisfy the condition

$$\frac{1}{\varepsilon(\delta_n)} \delta_n^2 \leq \alpha_n \alpha(\delta_n), \quad (4.31^*)$$

and such that

$$\| \tilde{A}X^{\alpha_n} - X^0 \| \geq \varepsilon_0 \quad (4.42)$$

for all numbers n .

Since the set $\{ X^{\alpha} \}$ is bounded, we can isolate a convergent subsequence from the sequence $\{ X^{\alpha_n} \}$ by virtue of the theorem of Bolzano-Weierstrass. In order not to change the designations, we assume that the whole sequence $\{ X^{\alpha_n} \}$ converges to

$$\text{a certain column } \bar{X}^0 = \begin{pmatrix} x_1^{-0} \\ \vdots \\ x_n^{-0} \end{pmatrix}, \text{ i.e. } \| X^{\alpha_n} - \bar{X}^0 \| \rightarrow 0 \text{ as } n \rightarrow \infty.$$

Let us make sure that

$$\| AX^{\alpha_n} - AX^0 \| \rightarrow 0 \text{ as } n \rightarrow \infty. \quad (4.43)$$

Indeed, using the triangle inequality, estimates (4.28), (4.29), (4.36) and (4.40) and relations (4.31*) we get

$$\begin{aligned} \| AX^{\alpha_n} - AX^0 \| &\leq \| AX^{\alpha_n} - \tilde{A}X^{\alpha_n} \| + \| \tilde{A}X^{\alpha_n} - \tilde{B}_{\tilde{A}} \| + \| \tilde{B}_{\tilde{A}} - AX^0 \| \\ &\leq \delta_n \| X^{\alpha_n} \| + \sqrt{F^{\alpha_n}(X^{\alpha_n}, \tilde{A}, \tilde{B}_{\tilde{A}})} + C\delta_n \leq \delta_n (\| X^{\alpha_n} \| + C) \\ &\quad + \sqrt{\alpha_n \varepsilon(\delta_n) C + \alpha_n \| X^0 \|^2} \rightarrow 0 \text{ as } n \rightarrow \infty. \end{aligned}$$

It immediately follows from inequality (4.43) that $A\bar{X}^0 = AX^0$, i.e. the limiting element of $\bar{X}^0$ is a solution of system (4.26) satisfying, by virtue of relation (4.41), the inequality $\| \bar{X}^0 \| \leq \| X^0 \|$. Since by definition the inverse inequality $\| X^0 \| \leq \| \bar{X}^0 \|$ holds true for the normal solution X^0 we have $\| \bar{X}^0 \| = \| X^0 \|$, i.e. $\bar{X}^0 = X^0$, but this contradicts inequality (4.42) which is valid for any number n .

The contradiction we have obtained complete the theorem proof of the theorem.

Chapter – 5

LINEAR OPERATORS

This chapter deals with the so-called **linear mappings** of linear and Euclidean spaces, i.e. mappings in which the image of the sum of elements is equal to the sum of their images and the image of the product of an element by a scalar is equal to the product of that scalar by the image of the element. We consider **complex** linear and Euclidean spaces here. The results referring to real spaces will be specified.

5.1. The Concept of a Linear Operator. Its Principal Properties

5.1.1. Definition of a linear operator. Let V and W be linear spaces whose dimensions are n and m respectively. We use the term **operator A from V into W** for the mapping of the form $A: V \rightarrow W$, associating every element x of the space V with a certain element y of the space W . In that case we use the designation $y = A(x)$ or $y = Ax$.

Definition. *The operator A from V into W is said to be linear if for any elements x_1 and x_2 of the space V and any complex number λ the following relations are satisfied:*

1°. $A(x_1 + x_2) = Ax_1 + Ax_2$ (the property of additivity of the operator).

2°. $A(\lambda x) = \lambda Ax$ (the property of homogeneity of the operator).

Remark 1. *If the space W is a complex plane, then the linear operator A from V into W is called a **linear form** or **linear functional**.*

Remark 2. If the space W coincides with the space V , then the linear operator acting in this case from V into V is also called a linear transformation of the space V .

5.1.2. Actions performed on linear operators. The space of linear operators. We shall define, in the set of all linear operators from V into W , the operations of **summation** of such operators and the **multiplication** of an operator by a scalar.

Suppose A and B are two linear operators from V into W . The sum of these operators is a linear operator $A + B$ defined by the equation

$$(A + B)x = Ax + Bx. \quad (5.1)$$

The **product** of the linear operator A by the scalar λ is a linear operator λA defined by the equation

$$(\lambda A)x = \lambda(Ax). \quad (5.2)$$

The **null** operator is an operator, designated as O , which maps all the elements of the space V into a null element of the space W .

To put it otherwise, the operator O acts according to the rule $Ox = 0$.

We define the **inverse operator** $-A$ of every operator A by the relation

$$-A = (-1)A.$$

It is easy to verify the validity of the following statement:

The set $L(V, W)$ of all linear operators from V into W , with the operations of summation and multiplication by a scalar indicated above and with a chosen null operator and its inverse forms a linear space.

5.1.3. The properties of the set $L(V, V)$ of linear operators. Let us discuss, in more detail, linear operators from V into V , i.e. investigate the set $L(V, V)$.

The **identity** (or **unit**) operator is a linear operator I acting according to the rule $Ix = x$ (here x is any element of V).

We now introduce the notion of the **product** of linear operators from the set $L(V, V)$.

The **product** of the operators A and B belonging to $L(V, V)$ is an operator AB acting according to the rule

$$(AB)x = A(Bx). \quad (5.3)$$

Note that in general $AB \neq BA$.

The following properties of linear operators from $L(V, V)$ hold true:

$$\left. \begin{aligned} 1^\circ. \lambda (AB) &= (\lambda A) B. \\ 2^\circ. (A + B) C &= AC + BC. \\ 3^\circ. A(B + C) &= AB + AC. \\ 4^\circ. (AB) C &= C + (BC). \end{aligned} \right\} \quad (5.4)$$

The first property (5.4) follows from the definition of the product of a linear operator by a scalar (see (5.2)) and the definition of the product of operators (see (5.3)).

Let us substantiate property 2°. In accordance with (5.1), (5.2) and (5.3), we have

$$\begin{aligned} ((A + B) C) x &= (A + B (Cx) = A (Cx) + B Cx) \\ &= (AC) x + (BC) x = (ACx + BC) x. \end{aligned}$$

Comparing the left-hand and the right-hand side of the last relations, we get an equality $(A + B) C = AC + BC$.

We have established property 2°.

Property 3° can be proved by a complete analogy.

Property 4° holds true since, in accordance with the definition (see (5.3)), the product of linear operators consists in their succession and, therefore, the linear operators $(AB)C$ and $A(BC)$ coincide and are consequently identical.

Remark 1. Property 4° makes it possible to define the product $AB...C$ of any finite number of operators from $L(V, V)$ and, in particular, the n th degree of the operator A by the formula

$$A^n = \underbrace{AA \dots A}_{n \text{ factors}}.$$

The relation $A^{n+m} = A^n A^m$ is obviously true.

We shall need the notion of an **inverse operator** of the given operator A from $L(V, V)$.

Definition 1. The linear operator B from $L(V, V)$ is said to be an inverse of the operator A from $L(V, V)$ if the relation $AB = BA = I$ is satisfied.

The inverse operator of the operator A is usually designated as A^{-1} .

It follows from the definition of the inverse operator A^{-1} that for any $x \in V$ there holds a relation $A^{-1}Ax = x$.

Thus, if $A^{-1}Ax = 0$, then $x = 0$, i.e. if the operator A has an inverse, then it follows from the condition $Ax = 0$ that $x = 0$.

We say that the linear operator A is one-to-one from V into V if any two distinct elements x_1 and x_2 are associated with different elements $y_1 = Ax_1$ and $y_2 = Ax_2$.

If the operator A is one-to-one from V into V , then the mapping $A: V \rightarrow V$ is a mapping of V onto V , i.e. every element $y \in V$ is an image of a certain element $x \in V$:

$$y = Ax.$$

To verify this fact, it is evidently sufficient to prove that n linearly independent elements $x_1, x_2, \dots, x_n$ of the space V are mapped by the operator A into n linearly independent elements $Ax_1, Ax_2, \dots, Ax_n$ of the same space.

Thus, assume that $x_1, x_2, \dots, x_n$ are linearly independent elements of V . If the linear combination $\alpha_1 Ax_1 + \alpha_2 Ax_2 + \dots + \alpha_n Ax_n$ is a zero element of the space V ,

$$\alpha_1 Ax_1 + \alpha_2 Ax_2 + \dots + \alpha_n Ax_n = 0.$$

then it follows from the definition of a linear operator (see 5.1.1.) that $A(\alpha_1 x_1 + \alpha_2 x_2 + \dots + \alpha_n x_n) = 0$.

Since the operator A is one-to-one from V into V , it follows from the last relation that $\alpha_1 x_1 + \alpha_2 x_2 + \dots + \alpha_n x_n = 0$. But since the elements $x_1, x_2, \dots, x_n$ are linearly independent, it follows that $\alpha_1 = \alpha_2 = \dots = \alpha_n = 0$. Consequently the elements $Ax_1, Ax_2, \dots, Ax_n$ are also linearly independent.

Note the following **assertion**.

For the linear operator A from $L(V, V)$ to have an inverse, it is necessary and sufficient that this operator be one-to-one from V into V .

Let us ascertain the necessity of this condition. Suppose the operator A has an inverse but is not one-to-one from V into V . This means that certain different elements x_1 and x_2 , $x_2 - x_1 \neq 0$ from V are associated with one and the same element $y = Ax_1 = Ax_2$. But then $A(x_2 - x_1) = 0$ and since A has an inverse, $x_2 - x_1 = 0$. But we pointed out above that $x_2 - x_1 \neq 0$. The contradiction obtained proves the necessity of the condition given in the assertion.

Let us now prove the *sufficiency* of the condition.

We assume that the operator A is one-to-one from V into V . Then every element $y \in V$ is associated with an element $x \in V$ such that $y = Ax$. There is, therefore, an operator A^{-1} possessing the property $A^{-1}y = A^{-1}(Ax) = x$. It is easy to verify that A^{-1} is a linear operator. By definition A^{-1} is an inverse of A . We have thus proved the sufficiency of the condition given in the assertion.

We introduce now the notions of the **kernel** and the **image** of a linear operator.

Definition 2. *The kernel of a linear operator A is the set of all the elements x of the space V for which $Ax = 0$.*

The kernel of the linear operator A is symbolized as $\ker A$.

If $\ker A = 0$, then the operator A is one-to-one from V into V . Indeed, in this case the condition $Ax = 0$ yields $x = 0$ and this means that different x_1 and x_2 are associated with different $y_1 = Ax_1$ and $y_2 = Ax_2$ (if y_1 were equal to y_2 , then $A(x_2 - x_1)$ would be zero, i.e. $x_1 = x_2$ and the elements x_1 and x_2 would not differ).

Thus, in accordance with the assertion proved above, *the condition $\ker A = 0$ is necessary and sufficient for the operator A to have an inverse.*

Definition 3. *The **image** of a linear operator A is the set of all elements y of the space V which can be represented as $y = Ax$. The image of the linear operator A is symbolized as $\operatorname{im} A$.*

Remark 2. Note that if $\ker A = 0$, then $\operatorname{im} A = V$ and vice versa. Therefore, alongside the condition $\ker A = 0$ indicated above, the condition $\operatorname{im} A = V$ is *also necessary and sufficient for the operator A to have an inverse.*

Remark 3. The kernel $\ker A$ and the image $\operatorname{im} A$ are evidently linear subspaces of the space V . We can, therefore, consider the *dimensions* $\dim(\ker A)$ and $\dim(\operatorname{im} A)$ of these subspaces.

The following theorem holds true.

Theorem 5.1. *Assume that the dimension $\dim V$ of the space V is equal to n and A is a linear operator from $L(V, V)$. Then*

$$\dim(\operatorname{im} A) + \dim(\ker A) = n.$$

Proof. Since $\ker A$ is a subspace of V , we can indicate a subspace V_1 of the space V such that V is a direct sum of V_1 and $\ker A$. In accordance with Theorem 2.10, $\dim V_1 + \dim(\ker A) = n$. Therefore, to prove the theorem, it is sufficient to make sure that $\dim V_1 = \dim(\operatorname{im} A)$.

Suppose $\dim V_1 = p$, $\dim(\operatorname{im} A) = q$ and $y_1, y_2, \dots, y_q$ is a basis in $\operatorname{im} A$. The linear operator A being one-to-one from V_1 into $\operatorname{im} A$, every element y from $\operatorname{im} A$ can be put into correspondence with a single element $x \in V_1$ such that $Ax = y$. Therefore, elements $x_1, x_2, \dots, x_q$ are defined in V_1 such that $Ax_k = y_k$, $k = 1, 2, \dots, q$. The elements $x_1, x_2, \dots, x_q$ are linearly independent since if $\alpha_1 x_1 + \alpha_2 x_2 + \dots + \alpha_q x_q = 0$, then $A(\alpha_1 x_1 + \alpha_2 x_2 + \dots + \alpha_q x_q) = \alpha_1 y_1 + \alpha_2 y_2 + \dots + \alpha_q y_q = 0$, and since the elements $y_1, y_2, \dots, y_q$ are linearly independent, we have $\alpha_1 = \alpha_2 = \dots = \alpha_q = 0$, i.e. $x_1, x_2, \dots, x_q$ are also linearly independent. Thus there are q linearly independent elements in V_1 . Consequently, $p \geq q$ (recall that $p = \dim V_1$).

Assume that $p > q$. Let us add elements $x_{q+1}, x_{q+2}, \dots, x_p$ to the linearly independent elements $x_1, x_2, \dots, x_q$ so that $x_1, x_2, \dots, x_p$ form a basis in V_1 . Since $p > q$ and $q = \dim(\operatorname{im} A)$, the elements $Ax_1, Ax_2, \dots, Ax_p$, belonging to $\operatorname{im} A$, are linearly dependent and, therefore, there exist numbers $\lambda_1, \lambda_2, \dots, \lambda_p$, not all equal to

zero, such that $\lambda_1 \mathbf{A} \mathbf{x}_1 + \lambda_2 \mathbf{A} \mathbf{x}_2 + \dots + \lambda_p \mathbf{A} \mathbf{x}_p = 0$. Hence $\mathbf{A}(\lambda_1 \mathbf{x}_1 + \lambda_2 \mathbf{x}_2 + \dots + \lambda_p \mathbf{x}_p) = 0$. Since $\mathbf{A}$ is one-to-one from V_1 into $\text{im } \mathbf{A}$, we get from the last equation $\lambda_1 \mathbf{x}_1 + \lambda_2 \mathbf{x}_2 + \dots + \lambda_p \mathbf{x}_p = 0$.

But $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p$ is a basis in V_1 . Therefore, $\lambda_1 = \lambda_2 = \dots = \lambda_p = 0$. It was indicated somewhat earlier that not all $\lambda_1, \lambda_2, \dots, \lambda_p$ are equal to zero. Consequently, the assumption $p > q$ leads to a contradiction. Thus $p = q$ and this completes the proof of the theorem.

The following theorem also holds true, which is reciprocal, in a certain sense, of Theorem 5.1.

Theorem 5.2. Suppose V_1 and V_2 are two subspaces of the n -dimensional space V such that $\dim V_1 + \dim V_2 = \dim V$. Then there is a linear operator $\mathbf{A}$ from $L(V, V)$ such that $V_1 = \text{im } \mathbf{A}$ and $V_2 = \ker \mathbf{A}$.

Proof. Suppose $\dim V_1 = p$ and $\dim V_2 = q$. We choose in the space V a basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ such that the elements $\mathbf{e}_{p+1}, \mathbf{e}_{p+2}, \dots, \mathbf{e}_n$ belong to V_2 . Next we take a basis $\mathbf{g}_1, \mathbf{g}_2, \dots, \mathbf{g}_p$ in space V_1 .

Let us now find the values of the linear operator $\mathbf{A}$ on the base vectors $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ of the space V as follows:

$$\mathbf{A} \mathbf{e}_1 = \mathbf{g}_1, \mathbf{A} \mathbf{e}_2 = \mathbf{g}_2, \dots, \mathbf{A} \mathbf{e}_p = \mathbf{g}_p,$$

$$\mathbf{A} \mathbf{e}_{p+1} = 0, \mathbf{A} \mathbf{e}_{p+2} = 0, \dots, \mathbf{A} \mathbf{e}_n = 0.$$

Furthermore, if $\mathbf{x} = x_1 \mathbf{e}_1 + x_2 \mathbf{e}_2 + \dots + x_p \mathbf{e}_p + x_{p+1} \mathbf{e}_{p+1} + \dots + x_n \mathbf{e}_n$, then $\mathbf{A} \mathbf{x} = x_1 \mathbf{g}_1 + x_2 \mathbf{g}_2 + \dots + x_p \mathbf{g}_p$. The operator $\mathbf{A}$ is, evidently linear and possesses the required properties. We have proved the theorem.

Let us introduce the notion of the **rank** of the linear operator $\mathbf{A}$.

The rank of the linear operator A is the number, symbolized as $\text{rank } A$ and equal to $\text{rank } A = \dim (\text{im } A)$.

Note the following obvious consequence of Theorem 5.1 and of Remark 2 in this subsection.

Corollary of Theorem 5.1. *For the operator A from $L(V, V)$ to have an inverse A^{-1} , it is necessary and sufficient that $\text{rank } A = \dim V = n$.*

Suppose A and B are linear operators from $L(V, V)$. The following theorem holds true.

Theorem 5.3. *The following relations hold true:*

$$\text{rank } AB \leq \text{rank } A, \text{rank } AB \leq \text{rank } B.$$

Proof. Let us prove the first of these relations. It is evident that $\text{im } AB \subseteq \text{im } A$. Therefore, $\dim (\text{im } AB) \leq \dim (\text{im } A)$, i.e. $\text{rank } AB \leq \text{rank } A$.

To prove the second relation, we make use of the following obvious inclusion: $\ker B \subseteq \ker AB$.

It follows from this inclusion that $\dim (\ker B) \leq \dim (\ker AB)$. The last inequality, in its turn, yields an inequality $\dim V - \dim (\ker AB) \leq \dim V - \dim (\ker B)$, and, in accordance with Theorem 5.1, we find from the last inequality that $\dim (\text{im } AB) \leq \dim (\text{im } B)$, i.e. $\text{rank } AB \leq \text{rank } B$. And this is what we had to prove.

Let us prove one more theorem on the ranks of linear operators.

Theorem 5.4. *Assume that A and B are linear operators from $L(V, V)$ and n is the dimension of V . Then $\text{rank } AB \geq \text{rank } A + \text{rank } B - n$.*

Proof. In accordance with Theorem 5.1.

$$\dim V = \dim (\text{im } AB) + \dim (\ker AB) = n. \quad (5.5)$$

Since $\text{rank } AB = \dim (\text{im } AB)$, we get from 5.5'

$$\text{rank } AB = n - \dim (\ker AB). \quad (5.6)$$

Since

$$\dim (\ker A) + \dim (\ker B) = 2n - (\text{rank } A + \text{rank } B) \quad (5.7)$$

in accordance with Theorem 5.1, it is sufficient to establish the inequality

$$\dim(\ker \mathbf{AB}) \leq \dim(\ker \mathbf{A}) + \dim(\ker \mathbf{B}) \quad (5.8)$$

in order to prove the theorem. Indeed, this inequality and relation (5.6) yield an inequality

$$\text{rank } \mathbf{AB} \geq n - (\dim(\ker \mathbf{A}) + \dim(\ker \mathbf{B}))$$

from which, in accordance with (5.7) immediately follows the validity of the assertion of the theorem.

We pass now to the substantiation of inequality (5.8).

Assume

$$\dim(\ker \mathbf{B}) = q. \quad (5.9)$$

According to Theorem 5.3, $\dim(\ker \mathbf{AB}) \geq q$. Therefore, the relation

$$\dim(\ker \mathbf{AB}) = p + q, \text{ where } p \geq 0 \quad (5.10)$$

holds true.

Since $\ker \mathbf{B} \subseteq \ker \mathbf{AB}$, we can choose a basis $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_{p+q}$ in the subspace $\ker \mathbf{AB}$ such that the elements $\mathbf{x}_{p+1}, \dots, \mathbf{x}_{p+q}$ form a basis in $\ker \mathbf{B}$. With such a choice of $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_{p+q}$ the elements $\mathbf{Bx}_1, \mathbf{Bx}_2, \dots, \mathbf{Bx}_p$ are linearly independent (if the linear combination

$$\sum_{k=1}^p \lambda_k = 0, \text{ then } \mathbf{B} \left(\sum_{k=1}^p \lambda_k \mathbf{x}_k \right) = 0, \text{ i.e. } \sum_{k=1}^p \lambda_k \mathbf{x}_k \subseteq \ker \mathbf{B} \text{ and this can}$$

happen, by virtue of the choice of $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p$, only for $(\lambda_k = 0, k = 1, 2, \dots, p)$.

Therefore, the elements $\mathbf{Bx}_1, \mathbf{Bx}_2, \dots, \mathbf{Bx}_p$ belong to $\ker \mathbf{A}$, i.e. $p \leq \dim(\ker \mathbf{A})$.

This inequality and relations (5.9) and (5.10) yield the required inequality (5.8). We have proved the theorem.

Corollary of Theorems 5.3 and 5.4. *If rank $\mathbf{A} = n$ (n being the dimension of V), then rank $\mathbf{AB} = \text{rank } \mathbf{BA} = \text{rank } \mathbf{B}$.*

This corollary follows from the inequalities

$\text{rank } \mathbf{AB} \leq \text{rank } \mathbf{B}$ (Theorem 5.3),

$\text{rank } \mathbf{AB} \geq \text{rank } \mathbf{B}$ (Theorem 5.4 for $\text{rank } \mathbf{A} = n$).

From these inequalities we find that $\text{rank } \mathbf{AB} = \text{rank } \mathbf{B}$.

The relation $\text{rank } \mathbf{BA} = \text{rank } \mathbf{B}$ can be proved by analogy.

5.2. Matrix Notation for Linear Operators

5.2.1. The matrices of linear operators in a given basis of the linear space V .

We fix the basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ in the linear space V .

Suppose $\mathbf{x}$ is an arbitrary element of V and

$$\mathbf{x} = \sum_{k=1}^n x^k \mathbf{e}_k \quad (5.11)$$

is the resolution of $\mathbf{x}$ into the given components.

Assume that $\mathbf{A}$ is a linear operator from $L(V, V)$. Then (5.11) yields

$$\mathbf{Ax} = \sum_{k=1}^n x^k \mathbf{Ae}_k. \quad (5.12)$$

Setting

$$\mathbf{Ae}_k = \sum_{j=1}^n a_k^j \mathbf{e}_j, \quad (5.13)$$

we rewrite (5.12) in the following form:

$$\mathbf{Ax} = \sum_{k=1}^n x^k \sum_{j=1}^n x_k^j \mathbf{e}_j = \sum_{j=1}^n \left(\sum_{k=1}^n a_k^j x^k \right) \mathbf{e}_j.$$

Thus, if $\mathbf{y} = \mathbf{Ax}$ and the element $\mathbf{y}$ has coordinates $y^1, y^2, \dots, y^n$, then

$$y^j = \sum_{k=1}^n a_k^j x^k, \quad j = 1, 2, \dots, n. \quad (5.14)$$

Let us consider a square matrix $\mathbf{A}$ with elements a_k^j :

$$A = (a_k^j).$$

This matrix is known as the **matrix of a linear operator in the given basis** $e_1, e_2, \dots, e_n$.

Besides the form of notation used for a linear operator indicated above, the matrix notation $y = Ax$ is also used for the given basis $e_1, e_2, \dots, e_n$ and if in this case $x = (x^1, x^2, \dots, x^n)$, then $y = (y^1, y^2, \dots, y^n)$, where $y^j, j = 1, 2, \dots, n$, are defined by means of relations (5.14) and the elements a_k^j , of the matrix A are calculated by formulas (5.13).

Remark 1. If the operator A is zero, then all the elements of the matrix A of that operator are equal to zero in any basis, i.e. A is a null matrix.

Remark 2. If A is an identity operator, i.e. $A = I$, then the matrix of that operator is an identity matrix in any basis. To put it otherwise, in this case $A = I$, where I is an identity matrix.

We have found that for the given basis of the linear space V , every linear operator A from $L(V, V)$ is associated with the matrix A of that operator. The reverse question naturally arises whether with a given basis in V every given matrix A can be put into correspondence with a linear operator A whose matrix is the given matrix. It is also important to find out whether the matrix of the linear operator in the given basis is unique.

The following statement holds true.

Theorem 5.5. *Suppose that a basis $e_1, e_2, \dots, e_n$ is given in the linear space V and that $A = (a_k^j)$ is a square matrix consisting of n rows and n columns. There is a unique linear operator A whose matrix in the given basis is the matrix*

Proof. Let us first prove the existence of the operator A . For that purpose we determine the values Ae_k of that operator on the base vectors e_k with the aid of relation (5.13), setting in this relation at, equal to the corresponding elements of the given matrix A . The value of the operator A on the arbitrary vector $x \subseteq V$, whose

representation in terms of the base vectors $e_1, e_2, \dots, e_n$ is given by formula (5.11), can be found by formula (5.12).

The operator we have constructed is, evidently; linear and the matrix A is a matrix of this operator.

The uniqueness of the operator A whose matrix in the basis $e_1, e_2, \dots, e_n$ is the matrix A follows from relations (5.13); these relations make it possible to define uniquely the values of the operator on the base vectors.

Remark 3. Assume that A and B are square matrices of order n , A and B are linear operators in the given basis $\{e_k\}$ of the space V corresponding to them. It follows from the proof of Theorem 5.5. that the matrix $A + \lambda B$, where λ is some number, is associated with a linear operator $A + \lambda B$ (recall that A , B , and $A + \lambda B$ belong to $L(V, V)$).

Let us prove the following theorem. _

Theorem 5.6. *The rank of the linear operator A is equal to the rank of the matrix A of that operator: $\text{rank } A = \text{rank } A$.*

Proof. By definition, $\text{rank } A = \dim(\text{im } A)$ and $\text{im } A$ is a linear closure of the vectors g_k :

$$g_k = \sum_{j=1}^n a_k^j e_j \quad (5.15)$$

(see the matrix form of notation for an operator and the definition of $\text{im } A$).

Therefore, $\text{rank } A$ is equal to the maximum number of linearly independent vectors g_k . Since the vectors $e_1, e_2, \dots, e_n$ are linearly independent, it follows, in accordance with (5.15), that the maximum number of linearly independent vectors g_k coincides with the maximum number of linearly independent rows $(a_k^1, a_k^2, \dots, a_k^n)$ of the matrix A , i.e. with the rank of A . We have proved the theorem.

Assume that A and B are arbitrary square matrices consisting of n rows and n columns. Theorems 5.3, 5.4, 5.5 and 5.6 have the following corollaries.

Corollary 1. The rank AB of the product of A and B satisfies the relations

$$\begin{aligned}\text{rank } AB &\leq \text{rank } A, \quad \text{rank } AB \leq \text{rank } B, \\ \text{rank } AB &\geq \text{rank } A + \text{rank } B - n.\end{aligned}$$

Corollary 2. The inverse operator A^{-1} of the operator A exists only in the case when the rank of the matrix A of the operator A is n ($n = \dim V$). Note that in this case there is also an inverse matrix A^{-1} of the matrix A .

5.2.2. Transformation of the matrix of a linear operator in passing to a new basis. Suppose V is a linear space, A is a linear operator from $L(V, V)$, $e_1, e_2, \dots, e_n$ and $\tilde{e}_1, \tilde{e}_2, \dots, \tilde{e}_n$ are two bases in V and

$$e_k = \sum_{i=1}^n u_k^i e_i, k = 1, 2, \dots, n \quad (5.16)$$

are formulas for passing from the basis $\{e_i\}$ to the basis $\{\tilde{e}_k\}$. Let us denote the matrix (u_k^i) by U :

$$U = (u_k^i). \quad (5.17)$$

Note that $\text{rank } U = n$. Suppose -

$$A = (a_k^j) \quad \text{and} \quad \tilde{A} = (\tilde{a}_k^j) \quad (5.18)$$

are the matrices of the operator A in the indicated bases. Let us find the connection between the matrices.

The following statement holds true.

Theorem 5.7. *The matrices A and $\tilde{A}$ of the operator A in the bases $\{e_k\}$ and $\{\tilde{e}_k\}$ respectively are connected by the relation*

$$A = U^{-1} \tilde{A} U.$$

where U^{-1} is an inverse of the matrix U defined by equation (5.17).

Proof. Directing our attention to the notion of the matrix of a linear operator, we get, according to (5.18),

$$A\mathbf{e}_k = \sum_{i=1}^n a_k^i \mathbf{e}_i, A\tilde{\mathbf{e}}_k = \sum_{i=1}^n \tilde{a}_k^i \tilde{\mathbf{e}}_i. \quad (5.19)$$

The definition of a linear operator, formulas (5.16) and the second formula, (5.19) yield relations .

$$A\tilde{\mathbf{e}}_k = A \left(\sum_{i=1}^n u_k^i \mathbf{e}_i \right), A\tilde{\mathbf{e}}_k = \sum_{j=1}^n \tilde{a}_k^j \sum_{i=1}^n u_i^j \mathbf{e}_j.$$

Therefore, the equality $\sum_{i=1}^n u_k^i A\mathbf{e}_i = \sum_{j=1}^n \left(\sum_{i=1}^n \tilde{a}_k^i u_i^j \right) \mathbf{e}_j$ holds true.

Substituting the expression $A\mathbf{e}_i$ into the left-hand side of this equation in accordance with the first formula (5.19), we find

$$\sum_{j=1}^n \left(\sum_{i=1}^n u_k^i a_i^j \right) \mathbf{e}_j, \sum_{j=1}^n \left(\sum_{i=1}^n \tilde{a}_k^i u_i^j \right) \mathbf{e}_j.$$

Since $\{ \mathbf{e}_j \}$ is a basis, the last relation yields equations

$$\sum_{i=1}^n u_k^i a_i^j = \sum_{i=1}^n \tilde{a}_k^i u_i^j, j, k = 1, 2, \dots, n. \quad (5.20)$$

If we direct our attention to the matrices A , $\tilde{A}$ and U (see (5.17) and (5.18)), we see that relations (5.20) are equivalent to the following matrix equation: $UA = \tilde{A}U$.

Multiplying both sides of this equation on the left by the matrix U^{-1} , we get the required relation $A = U^{-1} \tilde{A}U$.

And this is what we wished to prove.

Remark 1. Let us consider the formula $A = U^{-1} \tilde{A} U$. Multiplying both sides of the matrix equality on the left by the matrix U and on the right by U^{-1} , we obtain a relation

$$\tilde{A} = U A U^{-1} \quad (5.21)$$

which is another form of connection between the matrices A and $\tilde{A}$ of the linear operator A in different bases.

Remark 2. Suppose A and B are square matrices of order n , and A and B are linear operators in the given basis $\{e_k\}$ corresponding to them. As was pointed out (see Remark 3 in 5.2.1), the matrix $A + \lambda B$ is associated with a linear operator $A + \lambda B$. Let us find the form of the matrix of this operator in the basis $\{\tilde{e}_k\}$. Assume that $\tilde{A}$ and $\tilde{B}$ are the matrices of the operators A and B in the basis $\{\tilde{e}_k\}$. Then, in accordance with (5.21), we have

$$\tilde{A} = U A U^{-1}, \tilde{B} = U B U^{-1}. \quad (5.22)$$

In accordance with (5.21), the matrix of the linear operator $A + \lambda B$ in the basis $\{\tilde{e}_k\}$ has the form $U (A + \lambda B) U^{-1}$. Using the distributive property of matrix multiplication, we rewrite the last formula as follows (recall that this formula is a matrix of the linear operator $A + \lambda B$ in the basis $\{\tilde{e}_k\}$):

$$U A U^{-1} + \lambda (U B U^{-1}).$$

Considering relations (5.22), we see that the matrix of the operator $A + \lambda B$ in the basis $\{\tilde{e}_k\}$ is written as follows:

$$\tilde{A} + \lambda \tilde{B}.$$

In particular, if B is an identity matrix, $B = I$, then $\tilde{B} = I$ (see Remark 2 in 5.2.1 and Theorem 5.5) and, therefore, the matrix of the linear operator $A + \lambda I$ in the basis $\{\tilde{e}_k\}$ has the form $A + \lambda I$.

Corollary of Theorem 5.7. $\det A = \det \tilde{A} \dots$

In fact, the determinant of the product of matrices being equal to the product of the determinants of those matrices, it follows from the relation $A = U^{-1} \tilde{A} U$ that

$$\det A = \det U^{-1} \det \tilde{A} \det U. \quad (5.23)$$

Since $\det U^{-1} \det U = 1$, we obtain from relation (5.23) a relation $\det A = \det \tilde{A}$. Thus, the determinant of the matrix of a linear operator does not depend on the choice of a basis. We can, therefore, introduce the notion of the **determinant** $\det A$ of the linear operator A setting.

$$\det A = \det A, \quad (5.24)$$

where A is the matrix of the linear operator A in any basis.

5.2.3. A characteristic polynomial of a linear operator. Assume that A is a linear operator and I is an identity operator from $L(V, V)$.

Definition. A polynomial with respect to λ

$$\det (A - \lambda I) \quad (5.25)$$

is known as a **characteristic polynomial** of the operator A .

Suppose a basis $\{e_k\}$ is given in the basis V and $A = (a_k^j)$ is a matrix of the operator A in that basis. Then, in accordance with (5.24), the characteristic polynomial (5.25) of the operator A is

written as follows:

$$\det (A - \lambda I) = \begin{vmatrix} a_1^1 - \lambda & a_1^2 & \dots & a_1^n \\ a_2^1 & a_2^2 - \lambda & \dots & a_2^n \\ \dots & \dots & \dots & \dots \\ a_n^1 & a_n^2 & \dots & a_n^n - \lambda \end{vmatrix} \quad (5.26)$$

Let us write the characteristic polynomial (5.25) denoting, the coefficient in λ^k by d_k :

$$\det (A - \lambda I) = \sum_{k=0}^n d_k \lambda^k. \quad (5.21)$$

Remark 1. Since the value of the determinant $\det (A - \lambda I)$ does not depend on the choice of a basis, the coefficients d_k of the characteristic polynomial on the right-hand side of (5.27) do not depend on the choice of a basis either. Thus, *the coefficients d_k of the characteristic polynomial of the operator A are invariants*, i.e. quantities whose values are independent of the choice of a basis.

In particular, the coefficient d_{n-1} evidently equal to $a_1^1 + a_2^2 + \dots + a_n^n$ is an invariant. This invariant is known as a **trace of the operator A** and is designated as $\text{tr } A$:

$$\text{tr } A = a_1^1 + a_2^2 + \dots + a_n^n. \quad (5.28)$$

Remark 2. The equation _

$$\det (A - \lambda L) = 0 \quad (5.29)$$

is known as a **characteristic equation** of the operator A .

5.3. Eigenvalues and Eignevectors of Linear Operators

Suppose V_1 is a subspace of an n -dimensional linear space V , and A is a linear operator from $L(V, V)$.

Definition 1. *The space V_1 is known as an **invariant subspace** of the operator A if for every x belonging to V_1 the element Ax also belongs to V_1 .*

We can cite $\ker A$ and $\text{im } A$ as examples of invariant subspaces of the operator A .

Definition 2. *The number λ is called an **eigenvalue** of the operator A if there is a nonzero vector x such that*

$$Ax = \lambda x. \quad (5.30)$$

*In this case, the vector x is called an **eigenvector** of the operator A corresponding to the eigenvalue λ .*

The following statement holds true.

Theorem 5.8. *For the number λ to be an eigenvalue of the operator A , it is necessary and sufficient that this number should be a root of the characteristic equation (5.29) of the operator A .*

Proof. Suppose λ is an eigenvalue of the operator A , and $\mathbf{x}$ is an eigenvector corresponding to λ ($\mathbf{x} \neq 0$). Let us rewrite relation (5.30) in the following form:

$$(A - \lambda I) \mathbf{x} = 0.$$

Since $\mathbf{x}$ is a nonzero vector, it follows from the last equation that $\ker (A - \lambda I) \mathbf{x} \neq 0$, i.e.

$$\dim (\ker (A - \lambda I)) \geq 1. \quad (5.31)$$

Since in accordance with Theorem 5.1

$$\dim (\operatorname{im} (A - \lambda I)) + \dim (\ker (A - \lambda I)) = n,$$

we get from this equation and from inequality (5.31) that

$$\dim (\operatorname{im} (A - \lambda I)) \leq n - 1. \quad (5.32)$$

By definition, $\dim (\operatorname{im} (A - \lambda I))$ is equal to the rank of the operator $A - \lambda I$. Therefore, inequality (5.32) yields

$$\operatorname{rank} (A - \lambda I) < n. \quad (5.33)$$

Thus, if λ is an eigenvalue, then the rank of the matrix $A - \lambda I$ of the operator $A - \lambda I$ is smaller than n , i.e. $\det (A - \lambda I) = 0$ and, consequently, λ is a root of the characteristic equation.

Assume now that λ is a root of the characteristic equation (5.29). Then inequality (5.32) holds true and, consequently, inequality (5.31) also holds true, from which it follows that for the number λ there is nonzero vector $\mathbf{x}$ such that $(A - \lambda I) \mathbf{x} = 0$. The last relation is equivalent to relation (5.30). Therefore, λ is an eigenvalue. And this is what we had to prove.

Corollary. *Every linear operator has an eigenvalue.*

Indeed, a characteristic equation always possesses a root (by virtue of the fundamental theorem of algebra).

The following theorem holds true.

Theorem 5.9. *For the matrix A of the linear operator A in the given basis $\{e_k\}$ to be diagonal*, it is necessary and sufficient that the base vectors e_k be eigenvectors of that operator.*

Proof. Suppose the base vectors e_k are eigenvectors of the operator A . Then

$$Ae_k = \lambda_k e_k, \quad (5.34)$$

and, therefore, the matrix A of the operator A has the form (see relations (5.13), and the concept of the matrix of a linear operator)

$$A = \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_1 & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix} \quad (5.35)$$

i.e. is diagonal.

Suppose the matrix A of the linear operator A in; the given basis $\{e_k\}$ is diagonal, i.e. has the form (5.35). Then relations (5.13) assume the form (5.34), and this means that e_k are eigenvectors of the operator A . We have proved the theorem.

We shall prove one more property of eigenvectors.

Theorem 5.10. *Suppose the eigenvalues $\lambda_1, \lambda_2, \dots, \lambda_p$ of the operator A are distinct. Then the eigenvectors $e_1, e_2, \dots, e_p$ corresponding to them are linearly independent.*

Proof. We reason by induction. Since e_1 is a nonzero vector, the statement holds true for one vector ($p = 1$), one nonzero vector being linearly independent. Suppose we have proved the theorem for m vectors $e_1, e_2, \dots, e_m$. We add a vector e_{m+1} to these vectors and assume that there holds an equality

$$\sum_{k=1}^{m+1} \alpha_k e_k = 0. \quad (5.36)$$

Then, using the properties of linear operator, we get;

$$\sum_{k=1}^{m+1} \alpha_k A e_k = 0. \quad (5.37)$$

Since e_k are eigenvectors, we have $A e_k = \lambda_k e_k$ and, therefore, equation (5.37) can be rewritten as follows:

* Recall that a matrix is said to be diagonal if all its elements lying off the principal diagonal are equal to zero.

$$\sum_{k=1}^{m+1} \alpha_k \lambda_k e_k = 0. \quad (5.38)$$

In accordance with (5.36), $\sum_{k=1}^{m+1} \lambda_{m+1} \alpha_k e_k = 0$. Subtracting this equality from equality (5.38), we find that

$$\sum_{k=1}^m (\lambda_k - \lambda_{m+1}) \alpha_k e_k = 0. \quad (5.39)$$

All λ_k are distinct by the hypothesis, i.e. $\lambda_k - \lambda_{m+1} \neq 0$. Therefore it follows from (5.39) and from the assumption of a linear independence of the vectors $e_1, e_2, \dots, e_m$ that $\alpha_1 = \alpha_2 = \dots = \alpha_m = 0$. From this equality and from (5.36), and also from the condition that e_{m+1} is an eigenvector ($e_{m+1} \neq 0$) it follows that $\alpha_{m+1} = 0$. Thus, we get from equation (5.36) $\alpha_1 = \alpha_2 = \dots = \alpha_{m+1} = 0$. This means that the vectors $e_1, e_2, \dots, e_{m+1}$ are linearly independent. We have carried out induction and completed the proof of the theorem.

Corollary. *If the characteristic polynomial of the operator A has n distinct roots, then in some basis the matrix of the operator A has a diagonal form.*

Indeed, in the case being considered, in accordance with the theorem we have just proved, the eigenvectors are linearly independent and can, therefore, be chosen as

base vectors. But then, by Theorem 5.9, the matrix of the operator A is diagonal in that basis.

5.4. Linear and Sesquilinear Forms in a Euclidean Space

5.4.1. A special representation of a linear form in a Euclidean space. Suppose V is a Euclidean space, and C is a complex plane (one-dimensional complex linear space).

We introduced the notion of a **linear form**, i.e. a linear operator from V to C in 5.1.1. Here we shall get a special representation of an arbitrary linear form f from $L(V, C)$.

Lemma. Assume that f is a linear form from $L(V, C)$. Then there is a unique element h from V such that

$$f(\mathbf{x}) = (\mathbf{x}, \mathbf{h}). \quad (5.40)$$

Proof. To prove the existence of the element h , we choose an orthonormal basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ in V .

Let us consider the element h whose coordinates h^k in the chosen basis are defined by the relations

$$h^k = f(\overline{\mathbf{e}_k}). \quad (5.41)$$

Thus, $\mathbf{h} = \sum_{k=1}^n h^k \mathbf{e}_k$.

Assume that $\mathbf{x} = \sum_{k=1}^n x^k \mathbf{e}_k$ is an arbitrary element of the space V . Using the properties of the linear form f and equation (5.41), we obtain

$$f(\mathbf{x}) = \sum_{k=1}^n x^k f(\mathbf{e}_k) = \sum_{k=1}^n x^k \overline{h^k}. \quad (5.42)$$

Since in the orthonormal basis $\{\mathbf{e}_k\}$ the scalar product $(\mathbf{x}, \mathbf{h})$ of the vectors

$$\mathbf{x} = \sum_{k=1}^n x^k \mathbf{e}_k \quad \text{and} \quad \mathbf{h} = \sum_{k=1}^n h^k \mathbf{e}_k \quad \text{is equal to} \quad \sum_{k=1}^n x^k \overline{h^k}, \quad \text{we get}$$

$$f(\mathbf{x}) = (\mathbf{x}, \mathbf{h}) \quad \text{from (5.42).}$$

We have proved the existence of the vector $\mathbf{h}$.

Let us now prove the uniqueness of this vector. Assume that $\mathbf{h}_1$ and $\mathbf{h}_2$ are two vectors such that with their aid the form $f(\mathbf{x})$ can be represented as (5.40). It is evident that for any $\mathbf{x}$ there holds a relation $(\mathbf{x}, \mathbf{h}_1) = (\mathbf{x}, \mathbf{h}_2)$ which yields an equality $(\mathbf{x}, \mathbf{h}_2 - \mathbf{h}_1) = 0$. Setting $\mathbf{x} = \mathbf{h}_2 - \mathbf{h}_1$ in this equality and using the definition of the norm of an element in a Euclidean space, we find that $\|\mathbf{h}_2 - \mathbf{h}_1\| = 0$. Thus, $\mathbf{h}_2 = \mathbf{h}_1$ and this proves the lemma.

Remark. The lemma is evidently also valid when V is a real Euclidean space and $f \in L(V, R)$, where R is a real straight line.

5.4.2. Sesquilinear forms in a Euclidean space.: Special representation of such forms. We introduce the concept of a sesquilinear form in a linear space.

Definition. The number function $B(\mathbf{x}, \mathbf{y})$ whose arguments are various vectors $\mathbf{x}$ and $\mathbf{y}$ of the linear space L is known as a sesquilinear form if the relations

$$\left. \begin{aligned} B(\mathbf{x} + \mathbf{y}, \mathbf{z}) &= B(\mathbf{x}, \mathbf{z}) + B(\mathbf{y}, \mathbf{z}), \\ B(\mathbf{x}, \mathbf{y} + \mathbf{z}) &= B(\mathbf{x}, \mathbf{y}) + B(\mathbf{x}, \mathbf{z}), \\ B(\lambda \mathbf{x}, \mathbf{y}) &= \lambda B(\mathbf{x}, \mathbf{y}), \\ B(\mathbf{x}, \lambda \mathbf{y}) &= \overline{\lambda} B(\mathbf{x}, \mathbf{y}). \end{aligned} \right\} \quad (5.43)$$

hold for any vectors $\mathbf{x}, \mathbf{y}$ and $\mathbf{z}$ from L and any complex number λ .

In other words, the sesquilinear form $B(\mathbf{x}, \mathbf{y})$ is a number function of two vector arguments $\mathbf{x}, \mathbf{y}$ defined on various vectors $\mathbf{x}$ and $\mathbf{y}$ of the linear space L , which is linear with respect to the first argument $\mathbf{x}$ and antilinear with respect to the second argument $\mathbf{y}$.

Remark 1. If the linear space L is real, then the sesquilinear forms pass into the so-called bilinear forms, i.e. forms which are linear with respect to every argument

(since λ is real, the fourth relation (5.43) characterizes the linearity with respect to the second argument as well). Bilinear forms will be discussed in Chapter 7.

Let us consider a sesquilinear form given in a Euclidean space V . The following theorem on a special representation of such a form holds true.

Theorem 5.11. Suppose $B(x, y)$ is a sesquilinear form in the Euclidean space V . Then there is a unique linear operator A from $L(V, V)$ such that

$$B(x, y) = (x, Ay) \quad (5.44)$$

Proof. Assume that y is any fixed element of the space V . Then $B(x, y)$ is a linear form of the argument x . Therefore, in accordance with the lemma given in 5.4.1, there is a well-defined element h of the space V such that

$$B(x, y) = (x, h). \quad (5.45)$$

Thus, equation (5.45) puts every y from V into correspondence with a unique element h from V . We have thus defined the operator A such that $h = Ay$. The linearity of this operator follows directly from properties (5.43) of the sesquilinear form and from the properties of a scalar product.

Let us prove that the operator A is unique.

Suppose A_1 and A_2 are two operators such that with their aid the form $B(x, y)$ can be represented in the form (5.44). It is evident that the relation $(x, -A_1 y) = (x, -A_2 y)$ holds for any x and y , and it follows from the indicated relation that $(x, -A_2 y - A_1 y) = 0$. Setting $x = A_2 y - A_1 y$ in this equality and using the definition of the norm of an element we find that $\|A_2 y - A_1 y\| = 0$.

Thus, the equality $A_2 y = A_1 y$ is valid for any y from V , i.e. $A_1 = A_2$. We have proved the theorem.

Corollary. Suppose $B(x, y)$ is a sesquilinear form in the Euclidean space V . Then there is a unique linear operator A from $L(V, V)$ such that.

$$B(x, y) = (Ax, y). \quad (5.46)$$

The validity of this corollary follows from the following arguments. First, the form $B_1(x, y) = \overline{B(x, y)}$ is sesquilinear (this follows from the fact that $B(x, y)$ is a

sesquilinear form and from the definition of such a form). Furthermore, in accordance with Theorem 5.11 we get, for $B_1(x, y)$ a representation in the form

$$B_1(y, x) = (y, Ax). \quad (5.47)$$

Since the conjugate value of $\overline{B_1(x, y)}$ is equal to $B_1(x, y)$, we can take the conjugate values of the left-hand and right-hand sides of (5.47) and, bearing in mind that $B_1(x, y) = \overline{B(x, y)}$ we obtain

$$B(x, y) = \overline{B(y, Ax)}. \quad (5.48)$$

But $\overline{(y, Ax)} = (Ax, y)$ (see 4.3.1.). Therefore, (5.48) yields equation (5.46). We have proved the corollary.

Remark 2. Theorem 5.11 and the corollary of this theorem are also valid for the case when the Euclidean space is real. In that case, in the formulation of the theorem and the corollary the term “sesquilinear form” must be replaced by the term “bilinear form”. See also Remark 1.

Let us introduce the concept of a matrix of a sesquilinear form in the given basis $\{e_k\}$.

Assume that x and y belong to V , and $x = \sum_{j=1}^n x^j e_j$, $y = \sum_{k=1}^n y^k e_k$ are the resolution of x and y in terms of the base vectors $\{e_k\}$. The definition of a sesquilinear form yields relations

$$B\left(\sum_{j=1}^n x^j e_j, \sum_{k=1}^n y^k e_k\right) = \sum_{j=1}^n \sum_{k=1}^n x^j y^k B(e_j, e_k). \quad (5.49)$$

Setting

$$b_{jk} = B(e_j, e_k), \quad (5.50)$$

we write expression, (5.49) for $B(x, y)$ in the form

$$B(\mathbf{x}, \mathbf{y}) = \sum_{j=1}^n b_{jk} x^j \bar{y}^k.$$

The matrix $B = (b_{jk})$ is said to be the **matrix of the sesquilinear form** $B(\mathbf{x}, \mathbf{y})$ in the basis $\{\mathbf{e}_k\}$.

The following statement holds true.

Suppose that the sesquilinear form $B(\mathbf{x}, \mathbf{y})$ is represented in the form (5.46)

$$B(\mathbf{x}, \mathbf{y}) = (A\mathbf{x}, \mathbf{y}). \quad (5.46)$$

Suppose, furthermore, that the elements of the matrix A of the operator A in the given orthonormal basis are equal to a_j^k . Then in this basis $b_{jk} = a_j^k$.

To prove this statement, let us consider expression (5.50) for the coefficients b_{jk} in the sesquilinear form. We transform the right-hand side (5.50) with the aid of (5.46). In accordance with (5.13), we get

$$b_{jk} = B(\mathbf{e}_j, \mathbf{e}_k) = (A\mathbf{e}_j, \mathbf{e}_k) = \left(\sum_{q=1}^n a_j^q \mathbf{e}_q, \mathbf{e}_k \right) = \sum_{j=1}^n a_j^q (\mathbf{e}_q, \mathbf{e}_k).$$

The basis $\{\mathbf{e}_k\}$ being orthonormal, $(\mathbf{e}_q, \mathbf{e}_k) = 0$ if $q \neq k$ and $(\mathbf{e}_k, \mathbf{e}_k) = 1$. Therefore, from all the terms of the last sum, only that term is nonzero which results at $q = k$. Thus, $b_{jk} = a_j^k$. We have thus proved the statement.

Remark 3. If the sesquilinear form is represented as $B(\mathbf{x}, \mathbf{y}) = (\mathbf{x}, A\mathbf{y})$ and the elements of the matrix A of the operator A in the given orthonormal basis are equal to a_j^k , then $b_{jk} = a_j^{-k}$ in that basis.

5.5. Linear Self-Adjoint Operators in a Euclidean Space

5.5.1. The concept of an adjoint operator. We shall consider linear operators in a finite-dimensional Euclidean space V .

Definition 1. The operator A^* from $L(V, V)$ is said to be an adjoint or a conjugate operator of the linear operator A if the relation

$$(Ax, y) = (x, A^*y) \quad (5.51)$$

holds for any x and y from V .

It is easy to verify that the operator A^* , which is an adjoint of the linear operator A , is itself a linear operator. This follows from the obvious relation

$$\begin{aligned} (Ax, \alpha y_1 + \beta y_2) &= \overline{\alpha} (Ax, y_1) + \overline{\beta} (Ax, y_2) \\ &= \overline{\alpha} (x, A^* y_1) + \overline{\beta} (x, A^* y_2) = (x, A^*(\alpha y_1 + \beta y_2)), \end{aligned}$$

which is valid for any element x, y_1, y_2 and any complex numbers α and β .

Let us prove the following theorem.

Theorem 5.12. *Every linear operator A has a single adjoint.*

Proof. The scalar product (Ax, y) is evidently a sesquilinear form (see 4.3.1 and the definition of a sesquilinear form).

By Theorem 5.11, there is a unique linear operator A^* such that this form can be represented as (x, A^*y) . Thus, $(Ax, y) = (x, A^*y)$.

Consequently, the operator A^* is an adjoint of the operator A . The uniqueness of A^* follows from the uniqueness of the representation of a sesquilinear operator in the form (5.44). And this is what we wished to prove.

In what follows the symbol A^* designates an operator which is an adjoint of A .

Note the following properties of adjoint operators:

- 1°. $I^* = I$.
- 2°. $(A + B)^* = A^* + B^*$.
- 3°. $(\lambda A)^* = \overline{\lambda} A^*$.
- 4°. $(A^*)^* = A$.
- 5°. $(AB)^* = B^* A^*$.

It is easy to prove properties 1° - 4° and we invite the reader to do so. Here is the proof of property 5°.

In accordance with the definition of a product of operators, the relation $(AB)x = A(Bx)$ holds true. With the aid of this equality and the definition of an adjoint operator, we get the following chain of relations:

$$\begin{aligned} ((AB)x, y) &= (A(Bx), y) = (Bx, A^*y) \\ &= (x, B^*(A^*y)) = (x, (B^*A^*)y). \end{aligned}$$

Thus, $((AB)x, y) = (x, (B^*A^*)y)$. In other words, the operator B^*A^* is an adjoint of the operator AB . We have proved the validity of property 5°.

Remark. The notion of an adjoint operator for a real space can be introduced by analogy. The inferences made here and the properties of adjoint operators are also valid in this case. (here property 3° can be formulated as $(\lambda A)^* = \lambda A^*$).

5.5.2. Self-adjoint operators. Basic properties.

Definition 2. The linear operator A from $L(V, V)$ is *self-adjoint* if the relation

$$A^* = A$$

holds true. A self-adjoint operator in a real space can be defined by analogy.

An identity operator I can serve as a simple example of a self-adjoint operator (see property 1° of adjoint operators in 5.5.1).

Self-adjoint operators help in obtaining a special representation of arbitrary linear operators. Namely, the following statement holds true.

Theorem 5.13. Let A be a linear operator acting in a complex Euclidean space V . Then the representation $A = A_R + iA_I$ holds true, where A_R and A_I are self-adjoint operators called the real and the imaginary part of the operator A respectively.

Proof. In accordance with properties 2°, 3° and 4° of adjoint operators (see 5.5.1), the operators $A_R = \frac{(A + A^*)}{2}$ and $A_I = \frac{(A - A^*)}{2i}$ are self-adjoint.

Evidently, $A = A_R + iA_I$. We have proved the theorem.

In the following theorem we elucidate the conditions under which the products of self-adjoint operators are self-adjoint. We say that the operators A and B commute if $AB = BA$.

Theorem 5.14. *For the product $\mathbf{AB}$ of the self-adjoint operators $\mathbf{A}$ and $\mathbf{B}$ to be a self-adjoint operator, it is necessary and sufficient that they should commute.*

Proof. $\mathbf{A}$ and $\mathbf{B}$ being self-adjoint operators, the relations

$$(\mathbf{AB})^* = \mathbf{B}^* \mathbf{A}^* = \mathbf{BA} \quad (5.52)$$

are valid in accordance with property 5° of adjoint operators (see 5.5.1).

Consequently, if $\mathbf{AB} = \mathbf{BA}$, then $(\mathbf{AB})^* = \mathbf{AB}$, i.e. the operator $\mathbf{AB}$ is self-adjoint. Now if $\mathbf{AB}$ is a self-adjoint operator, then $\mathbf{AB} = (\mathbf{AB})^*$, and then, by virtue of (5.52), $\mathbf{AB} = \mathbf{BA}$. And this is what we set out to prove.

In the theorems that follow we establish a number of significant properties of self-adjoint operators.

Theorem 5.15. *If the operator $\mathbf{A}$ is self-adjoint, then for any $\mathbf{x} \in V$ the scalar product $(\mathbf{Ax}, \mathbf{x})$ is a real number.*

Proof. The validity of the statement of the theorem follows from the property of a scalar product in a complex Euclidean space $(\mathbf{Ax}, \mathbf{x}) = \overline{(\mathbf{x}, \mathbf{Ax})}$ and the definition of a self-adjoint operator $(\mathbf{Ax}, \mathbf{x}) = (\mathbf{x}, \mathbf{Ax})$.

Theorem 5.16. *The eigenvalues of a self-adjoint operator are real.*

Proof. Let λ be an eigenvalue of the self-adjoint operator $\mathbf{A}$. By definition of an eigenvalue of the operator $\mathbf{A}$ (see Definition 2 in 5.3), there is a nonzero vector $\mathbf{x}$ such that $\mathbf{Ax} = \lambda \mathbf{x}$. It follows from this relation that the real (by virtue of Theorem 5.15) scalar product $(\mathbf{Ax}, \mathbf{x})$ can be represented as

$$(\mathbf{Ax}, \mathbf{x}) = \lambda (\mathbf{x}, \mathbf{x}) = \lambda \|\mathbf{x}\|^2.$$

Since $\|\mathbf{x}\|^2$ and $(\mathbf{Ax}, \mathbf{x})$ are real, λ is evidently also real. We have proved the theorem.

In the following theorem we elucidate the property of orthogonality of the eigenvectors of a self-adjoint operator.

Theorem 5.17. *If $\mathbf{A}$ is a self-adjoint operator, then the eigenvectors corresponding to various eigenvalues of that operator are orthogonal*

Proof. Let λ_1 and λ_2 be various eigenvalues ($\lambda_1 \neq \lambda_2$) of the self-adjoint operator A , and $\mathbf{x}_1$ and $\mathbf{x}_2$ be the eigenvectors corresponding to them. Then there hold relations $A\mathbf{x}_1 = \lambda_1\mathbf{x}_1$, $A\mathbf{x}_2 = \lambda_2\mathbf{x}_2$. Therefore, the scalar products $(A\mathbf{x}_1, \mathbf{x}_2)$ and $(\mathbf{x}_1, A\mathbf{x}_2)$ are respectively equal to the expressions

$$(A\mathbf{x}_1, \mathbf{x}_2) = \lambda_1 (\mathbf{x}_1, \mathbf{x}_2), (\mathbf{x}_1, A\mathbf{x}_2) = \lambda_2 (\mathbf{x}_1, \mathbf{x}_2).$$

The operator A being self-adjoint, the scalar products $(A\mathbf{x}_1, \mathbf{x}_2)$ and $(\mathbf{x}_1, A\mathbf{x}_2)$ are equal and, therefore, by means of subtraction we get from the last relations an equation

$$(\lambda_2 - \lambda_1) (\mathbf{x}_1, \mathbf{x}_2) = 0.$$

Since $\lambda_2 \neq \lambda_1$, it follows from the last equation that the scalar product $(\mathbf{x}_1, \mathbf{x}_2)$ is equal to zero, i.e. that the eigenvectors $\mathbf{x}_1$ and $\mathbf{x}_2$ are orthogonal. And this is what we wished to prove.

5.5.3. The norm of a linear operator. Let A be a linear operator mapping the Euclidean space V into the same space. We introduce the concept of the norm of the operator A .

Definition 3. The *norm* $\|A\|$ of the linear operator A is a number defined by the relation

$$\|A\| = \sup_{\|\mathbf{x}\|=1} \|A\mathbf{x}\|. \quad (5.53)$$

The definition of the norm of a linear operator yields the following obvious inequality:

$$\|A\mathbf{x}\| \leq \|A\| \|\mathbf{x}\| \quad (5.54)$$

(to prove it, it is sufficient to use the relation $A\mathbf{x} = \left(A \frac{\mathbf{x}}{\|\mathbf{x}\|}\right) \|\mathbf{x}\|$). It follows from relations (5.54) that if $\|A\| = 0$, then the operator A is zero.

We can define the norm of the self-adjoint operator A in another

way. Namely, the following statement holds true. ;

If A is a self-adjoint operator, then the norm $\|A\|$ of the operator A introduced above is ' $\sup_{\|x\|=1} |(Ax, x)|$:

$$\sup_{\|x\|=1} |(Ax, x)| = \|A\|. \quad (5.55)$$

Proof. The Cauchy-Bunyakovsky inequality (see 4.3.2)

$$|(Ax, x)| \leq \|Ax\| \|x\|$$

is valid for any x from V . From this inequality and inequality (5.54) we get an inequality $|(Ax, x)| \leq \|Ax\| \|x\|^2$. Therefore, the number μ

$$\mu = \sup_{\|x\|=1} |(Ax, x)| \quad (5.56)$$

satisfies the relation

$$\mu \leq \|A\|. \quad (5.57)$$

Note that the equation $((Az, z) = \left(A \frac{z}{\|z\|}, \frac{z}{\|z\|}\right) \|z\|^2$, $z \neq 0$, and the definition of the number μ (see (5.56)) yields an inequality

$$|(Az, z)| \mu \leq \|z\|^2.$$

Let us draw attention to the obvious identity

$$4 \operatorname{Re} (Ax, y) = (A(x + y), x + y) - (A(x - y), x - y)$$

(the symbol $\operatorname{Re} (Ax, y)$ in this identity is used for the real part of the complex number (Ax, y) ; the identity itself immediately follows from the properties of a scalar product, see 4.3.1.). Taking the absolute values of the expressions on the left-hand and right-hand sides of this identity, using the property of the modulus of a sum and inequality (5.58), we get the following relations*:

$$4 |\operatorname{Re} (Ax, y)| \leq \|x + y\|^2 + \mu \|x - y\|^2 = 2\mu (\|x\|^2 + \|y\|^2).$$

From this, for $\|x\| = \|y\| = 1$, we get an inequality

$$|\operatorname{Re}(\mathbf{Ax}, \mathbf{y})| \leq.$$

Setting $\mathbf{y} = \mathbf{Ax} / \|\mathbf{Ax}\|$ in this inequality (it is evident that $\|\mathbf{y}\| = 1$) and bearing in mind that the number $(\mathbf{Ax}, \mathbf{Ax}) = \|\mathbf{Ax}\|^2$ is real (and consequently, $\operatorname{Re}(\mathbf{Ax}, \mathbf{Ax}) = (\mathbf{Ax}, \mathbf{Ax}) = \|\mathbf{Ax}\|^2$), we obtain $\|\mathbf{Ax}\| \leq \mu, \|\mathbf{x}\| = 1$. From this, in accordance with inequality (5.53), we find that $\|\mathbf{Ax}\| \leq \mu$.

To complete the proof, we must compare the inequality obtained with inequality (5.57) and use the definition of the number μ (see (5.56)).

5.5.4. Other properties of self-adjoint operators. We shall prove here a number of significant properties of linear operators connected with the concept of a norm. We shall first establish the necessary and sufficient condition for an operator to be self-adjoint. Let us prove the following theorem.

Theorem 5.18. *For the linear operator $\mathbf{A}$ to be self-adjoint, it is necessary and sufficient that*

$$\operatorname{Im}(\mathbf{Ax}, \mathbf{x}) = 0^{**}.$$

Proof. By Theorem 5.13, the arbitrary linear operator $\mathbf{A}$ can be represented as $\mathbf{A} = \mathbf{A}_R + i\mathbf{A}_I$, where $\mathbf{A}_R$ and $\mathbf{A}_I$ are self-adjoint operators. Therefore,

$$(\mathbf{Ax}, \mathbf{x}) = (\mathbf{A}_R \mathbf{x}, \mathbf{x}) + i(\mathbf{A}_I \mathbf{x}, \mathbf{x}),$$

and, in accordance with Theorem (5.15), the numbers $(\mathbf{A}_R \mathbf{x}, \mathbf{x})$ and $(\mathbf{A}_I \mathbf{x}, \mathbf{x})$ are real for any $\mathbf{x}$. Consequently, these numbers are equal, respectively, to the real and the imaginary part of the complex number $(\mathbf{Ax}, \mathbf{x})$:

$$\operatorname{Re}(\mathbf{Ax}, \mathbf{x}) = (\mathbf{A}_R \mathbf{x}, \mathbf{x}), \operatorname{Im}(\mathbf{Ax}, \mathbf{x}) = (\mathbf{A}_I \mathbf{x}, \mathbf{x}).$$

We assume that $\mathbf{A}$ is a self-adjoint operator.

By Theorem 5.15, in this case the number $(\mathbf{Ax}, \mathbf{x})$ is real and, therefore, $\operatorname{Im}(\mathbf{Ax}, \mathbf{x}) = 0$. We have proved the necessity of the hypothesis.

Let us now prove its sufficiency.

Suppose $\operatorname{Im}(\mathbf{Ax}, \mathbf{x}) = (\mathbf{A}_I \mathbf{x}, \mathbf{x}) = 0$. Hence $\|\mathbf{A}_I\| = 0$, i.e. $\mathbf{A}_I = 0$. Therefore, $\mathbf{A} = \mathbf{A}_R$, where $\mathbf{A}_R$ is a self-adjoint operator. We have thus proved the theorem.

In the following assertions we elucidate certain properties of the eigenvalues of self-adjoint operators.

* We have used the definition of the norm of an element in a complex Euclidean space.

**The symbol $\text{Im} (\mathbf{Ax}, \mathbf{x})$ designates the imaginary part of the complex number $(\mathbf{Ax}, \mathbf{x})$. The equality $\text{Im} (\mathbf{Ax}, \mathbf{x}) = 0$ means that the number $(\mathbf{Ax}, \mathbf{x})$ is real.

Lemma. *Any eigenvalue λ of the arbitrary linear self-adjoint operator $\mathbf{A}$ in a Euclidean space is equal to the scalar product $(\mathbf{Ax}, \mathbf{x})$, where $\mathbf{x}$ is some vector satisfying the condition $\|\mathbf{x}\| = 1$:*

$$\lambda = (\mathbf{Ax}, \mathbf{x}), \|\mathbf{x}\| = 1.$$

Proof. Since λ is an eigenvalue of the operator $\mathbf{A}$, there is a non-zero vector $\mathbf{z}$ such that

$$\mathbf{Az} = \lambda \mathbf{z}. \quad (5.60)$$

Setting $\mathbf{x} = \mathbf{z} / \|\mathbf{z}\|$ (it is evident that $\|\mathbf{x}\| = 1$), we rewrite (5.60) as follows: $\mathbf{Ax} = \lambda \mathbf{x}$, $\|\mathbf{x}\| = 1$. Hence we get a relation $(\mathbf{Ax}, \mathbf{x}) = \lambda (\mathbf{x}, \mathbf{x}) = \lambda \|\mathbf{x}\|^2 = \lambda$, i.e. (5.59) holds true. We have proved the lemma.

Corollary. *Suppose $\mathbf{A}$ is a self-adjoint operator and λ is any eigenvalue of that operator. Suppose furthermore that*

$$m = \inf_{\|\mathbf{x}\|=1} (\mathbf{Ax}, \mathbf{x}), M = \sup_{\|\mathbf{x}\|=1} (\mathbf{Ax}, \mathbf{x}). \quad (5.61)$$

The following inequalities hold true:

$$m \leq \lambda \leq M. \quad (5.62)$$

Remark 1. Since the scalar product $(\mathbf{Ax}, \mathbf{x})$ is a continuous function of $\mathbf{x}$, this function is bounded on the closed set $\|\mathbf{x}\| = 1$ and attains its least upper bound M and greatest lower bound m (called a supremum and an infimum respectively).

Remark 2. In accordance with Theorem 5.16, the eigenvalues of a self-adjoint operator are real. Inequalities (5.62) are, therefore, meaningful.

Proof. of the corollary. Since any eigenvalue λ satisfies relation (5.59), every eigenvalue is, evidently, confined between the least up-per bound M and the greatest lower bound m of the scalar product $(\mathbf{Ax}, \mathbf{x})$. inequalities (5.62) are, therefore, valid.

We shall prove that the numbers m and M defined by relations (5.16) are, respectively, the least and the greatest value of the self-adjoint operator $\mathbf{A}$. We shall first verify the validity of the following statement.

5.19. *Assume that $\mathbf{A}$ is a self-adjoint operator and addition, $(\mathbf{Ax}, \mathbf{x}) \geq 0$ for any $\mathbf{x}$. Then $\|\mathbf{A}\|$ is equal to the greatest eigenvalue of that operator.*

Proof. We have remarked (see the statement made in 5.5.3) that $\|\mathbf{A}\| = \sup_{\|\mathbf{x}\|=1} (\mathbf{Ax}, \mathbf{x})$. Since $(\mathbf{Ax}, \mathbf{x}) \geq 0$, we have $\|\mathbf{A}\| = \sup_{\|\mathbf{x}\|=1} (\mathbf{Ax}, \mathbf{x})$.

According to Remark 1 of this subsection,

$$(\mathbf{Ax}_0, \mathbf{x}_0) = \|\mathbf{A}\| \lambda$$

for a certain $\mathbf{Ax}_0, \|\mathbf{x}_0\| = 1$.

Using the definition of a norm and the equations we have just written, we obtain relations.

$$\begin{aligned} \|(\mathbf{A} - \lambda \mathbf{I}) \mathbf{x}_0 \|^2 &= \| \mathbf{Ax}_0 \|^2 - 2\lambda (\mathbf{Ax}_0, \mathbf{x}_0) + \lambda^2 \|\mathbf{x}_0\|^2 \\ &= \|\mathbf{A}\|^2 - 2\|\mathbf{A}\| \cdot \|\mathbf{A}\| + \|\mathbf{A}\|^2 \cdot 1 = 0. \end{aligned}$$

Thus, $(\mathbf{A} - \lambda \mathbf{I}) \mathbf{x}_0 = 0$, or, to put it otherwise, $\mathbf{Ax}_0, \lambda \mathbf{x}_0$, is an eigenvalue of the operator $\mathbf{A}$. The fact that λ is the greatest eigenvalue follows from the corollary of the last lemma. We have thus proved the theorem.

Let us now prove that the numbers m and M (see (5.61)) are, respectively, the least and the greatest eigenvalue of the self-adjoint operator $\mathbf{A}$.

Theorem 5.20. *Let $\mathbf{A}$ be a self-adjoint operator, and m and M are, respectively, the greatest lower and the least upper bound of $(\mathbf{Ax}, \mathbf{x})$ on the set $\|\mathbf{x}\| = 1$. These numbers are the least and the greatest eigenvalue of the operator $\mathbf{A}$.*

Proof. It is, evidently, sufficient to prove that the numbers m and M are eigenvalues of the operator A . Then, it follows immediately from inequalities (5.62) that m and M are, respectively, the least and the greatest eigenvalue.

We shall first prove that m is an eigenvalue. For that purpose, we shall consider a self-adjoint operator $B = A - mI$. Since

$$(Bx, x) = (Ax, x) - m(x, x) \geq 0.$$

the operator B satisfies the conditions of Theorem 5.19 and, therefore, the norm $\|B\|$ of that operator is equal to its greatest eigenvalue. On the other hand,

$$\|B\| = \sup_{\|x\|=1} (Bx, x) = \sup_{\|x\|=1} (Ax, x) - m = M - m.$$

Thus, $(M - m)$ is the greatest eigenvalue of the operator B . Consequently, there is nonzero vector x_0 such that

$$Bx_0 = (M - m)x_0. \quad (5.63)$$

Since $B = A - mI$, we have $Bx_0 = Ax_0 - mx_0$. Substituting the expression Bx_0 into left-hand side of (5.63), we get, after simple transformations, a relation $Ax_0 = Mx_0$. Thus, M is an eigenvalue of the operator A .

Let us now make sure that the number m is also an eigenvalue of the operator A .

We shall consider the self-adjoint operator $B = -A$. Evidently, $-m = \sup_{\|x\|=1} (Bx, x)$. According to the proof we have just carried out, the number

$-m$ is an eigenvalue of the operator B . Since $B = -A$, m is an eigenvalue of the operator A . We have proved the theorem.

In the following theorem we elucidate a significant property of the eigenvectors of a self-adjoint operator.

Theorem 5.21. *Every self-adjoint linear operator A acting in an n -dimensional Euclidean space V possesses n linearly independent pairwise orthogonal unit eigenvectors.*

Proof. Let λ_1 be the maximum eigenvalue of the operator A ($\lambda_1 = \sup_{\|x\|=1} (Ax, x)$).

We denote by e_1 the eigenvector corresponding to λ_1 and satisfying the condition $\|e_1\|=1$ (the possibility of choosing it follows from the proof of the lemma in this subsection).

We denote by V_1 an $(n - 1)$ dimensional subspace of the space V which is orthogonal to e_1 . Evidently, V_1 is an invariant subspace of the operator A (i.e. if $e \in V_1$ then $Ax \in V_1$ as well). Indeed, suppose $x \in V_1$ (i.e. $(x, e_1) = 0$). Then*

$$(Ax, e_1) = (x, Ae_1) = \lambda_1 (x, e_1) = 0.$$

Consequently, Ax is an element of V_1 and, therefore, V_1 is an invariant subspace of the operator A . This enables us to consider the operator A in the subspace V_1 . In this subspace A is a self-adjoint operator. Consequently, there is a maximum eigenvalue λ_2 of that operator which can be found by the relation',

$$\lambda_2 = \max_{\substack{\|x\|=1 \\ x \perp e_1}} (Ax, x).$$

In addition, there is a vector $e_2, e_2 \perp e_1, \|e_2\|=1$ such that $Ae_2 = \lambda_2 e_2$.

Considering then an $(n - 2)$ dimensional subspace V_2 which is orthogonal to the vectors e_1 and e_2 and repeating the arguments presented above, we construct an eigenvector $e_3, \|e_3\|=1$ which is orthogonal to e_1 and e_2 . Continuing with the same arguments, we successively find n mutually orthogonal eigenvectors $e_1, e_2, \dots, e_n$ satisfying the condition $\|e_i\|=1, i = 2, \dots, n$.

Remark 1. In what follows we agree to enumerate the eigenvalues of a self-adjoint operator in a decreasing order, with the repeating, i.e. multiple, eigenvalues taken into account. In this case, $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$ and the eigenvectors $e_1, e_2, \dots, e_n$

corresponding to them can be considered mutually orthogonal and satisfying the condition $\|e_i\| = 1$. Thus,

$$(e_i, e_j) = \begin{cases} 1 & \text{for } i = j, \\ 0 & \text{for } i \neq j. \end{cases}$$

Remark 2. It follows from the arguments presented in the proof of Theorem 5.21 that $\lambda_{m+1} = \max_{\substack{x \perp e_k \\ k=1, 2, \dots, m}} \frac{(Ax, x)}{(x, x)}$. This relation can also be written in the form

$\lambda_{m+1} = \max_{x \perp E_m} \frac{(Ax, x)}{(x, x)}$, where E_m is a linear closure of the vectors $e_1, e_2, \dots, e_m$. The validity of the remark follows from the fact that $(x, x) = \|x\|^2$ and, therefore

$$\frac{(Ax, x)}{(x, x)} = \left(A \frac{x}{\|x\|}, \frac{x}{\|x\|} \right),$$

with the norm of the element $x / \|x\|$ being equal to unity.

Let ξ_m be the set of all m -dimensional subspaces of the space V . The following significant **minimax** property of eigenvalues holds true.

Theorem 5.22. Suppose A is a self-adjoint operator and $\lambda_1, \lambda_2, \dots, \lambda_n$ are its eigenvalues enumerated in the 'order indicated in Remark 1. Then

$$\lambda_{m+1} = \min_{E \in \xi_m} \max_{x \perp E} \frac{(Ax, x)}{(x, x)}. \quad (5.64)$$

Proof. Suppose E_m is a linear closure of the eigenvectors $e_1, e_2, \dots, e_m$ of the operator A (see Remark 1). By virtue of Remark 2

$$\max_{x \perp E_m} \frac{(Ax, x)}{(x, x)} = \lambda_{m+1}.$$

Therefore, to prove the theorem, it is sufficient to verify the validity of the relation

$$\max_{x \perp E \subset \xi_m} \frac{(Ax, x)}{(x, x)} \geq \max_{x \in E_m} \frac{(Ax, x)}{(x, x)} = \lambda_{m+1} \quad (5.65)$$

for any $E \in \xi_m$.

Let us prove relation (5.65).

We denote by $E^\perp$ the orthogonal complement of the space E (see 4.2.3). It follows from Theorem 2.10 that the dimension of $E^\perp$ is $n - m$. Consequently,

$$\dim E^\perp + \dim E_{m+1} = (n - m) + (m + 1) = n + 1 > n.$$

By Theorem 2.9 this means that the intersection of the subspaces $E^\perp$ and E_{m+1} contains a nonzero element. Thus, there exists an element $\tilde{x}$ such that

$$\tilde{x} \perp E, \|\tilde{x}\| = 1, \tilde{x} \in E_{m+1}, \text{ i.e. } \tilde{x} = \sum_{k=1}^{m+1} c_k \mathbf{e}_k. \quad \text{Since } \|\tilde{x}\| = 1 \text{ and the basis}$$

$\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_{m+1}$ is orthonormal it follows, by virtue of the Pythagoras Theorem (see 4.1.2), that

$$\|\tilde{x}\|^2 = \sum_{k=1}^{m+1} |c_k|^2 = 1. \quad (5.66)$$

We also have $A\tilde{x} = \sum_{k=1}^{m+1} c_k A\mathbf{e}_k = \sum_{k=1}^{m+1} c_k \lambda_k \mathbf{e}_k$. Since $\mathbf{e}_k$, are eigenvectors of the

operator A , the last relations yield $A\tilde{x} = \sum_{k=1}^{m+1} c_k \lambda_k \mathbf{e}_k$. From this and from the

orthonormality of $\mathbf{e}_k$ we find that the relation

$$(A\tilde{x}, x) = \left(\sum_{k=1}^{m+1} c_k \lambda_k \mathbf{e}_k, \sum_{p=1}^{m+1} c_p \mathbf{e}_p \right) = \sum_{k=1}^{m+1} |c_k|^2 \lambda_k A\mathbf{e}_k. \quad (5.67)$$

holds true.

We have enumerated the eigenvalues in a decreasing order with their possible multiplicity taken into account. Therefore, $\lambda_{m+1} \geq \lambda_k, k = 1, 2, \dots, m$. From this and from (5.67) and (5.66)

$$(A\mathbf{x}, \mathbf{x}) = \sum_{k=1}^{m+1} |c_k|^2 \lambda_k \geq \lambda_{m+1} \sum_{k=1}^{m+1} |c_k|^2 = \lambda_{m+1}.$$

Noting that the norm of the element $\mathbf{x} / \|\mathbf{x}\|$ is equal to 1 and if $\|\tilde{\mathbf{x}}\| = 1$ for any $\mathbf{x} \neq 0$ and taking into account that $\mathbf{x} \perp E$, we get

$$\max_{\mathbf{x} \perp E} \frac{(A\mathbf{x}, \mathbf{x})}{(\mathbf{x}, \mathbf{x})} = \max_{\mathbf{x} \perp E} \left(A \frac{\mathbf{x}}{\|\mathbf{x}\|}, \frac{\mathbf{x}}{\|\mathbf{x}\|} \right) \geq (A\tilde{\mathbf{x}}, \tilde{\mathbf{x}}) \geq \lambda_{m+1}.$$

We have thus established relations (5.65) and proved the theorem.

5.5.5. Expansion of self-adjoint operators in terms of eigenfunctions. Hamilton-Cayley theorem. Let us consider a self-adjoint operator A and its eigenvalues $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$. In this case $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ is an orthonormal basis consisting of eigenvectors corresponding to $\{\lambda_i\}$. Let $\mathbf{x} \in V$. Then

$$\mathbf{x} = \sum_{k=1}^n (\mathbf{x}, \mathbf{e}_k) \mathbf{e}_k \quad (5.68)$$

(see 4.2.3), and since $A\mathbf{e}_k = \lambda_k \mathbf{e}_k$, we get, by means of (5.68).

$$A\mathbf{x} = \sum_{k=1}^n \lambda_k (\mathbf{x}, \mathbf{e}_k) \mathbf{e}_k. \quad (5.69)$$

The operator P_k , defined by the relation

$$P_k \mathbf{x} = (\mathbf{x}, \mathbf{e}_k) \mathbf{e}_k, \quad (5.70)$$

is known as a **projection** onto the one-dimensional subspace generated by the vectors $\mathbf{e}_k$.

It immediately follows from the properties of a scalar product that $\mathbf{P}_k$ is a self-adjoint linear operator.

Note the following significant properties of projections:

1°. $\mathbf{P}_k^2 = \mathbf{P}_k$ (hence $\mathbf{P}_k^m = \mathbf{P}_k$, where m is natural).

2°. $\mathbf{P}_k \mathbf{P}_j = 0$, where $k \neq j$.

The proof of these properties follows from the relations

$$\begin{aligned} (\mathbf{P}_k \mathbf{P}_j) \mathbf{x} &= \mathbf{P}_k (\mathbf{P}_j \mathbf{x}) = \mathbf{P}_k (\mathbf{x}, \mathbf{e}_j) \mathbf{e}_j \\ &= (\mathbf{x}, \mathbf{e}_j) \mathbf{e}_j (\mathbf{e}_j, \mathbf{e}_k) \mathbf{e}_k = \begin{cases} (\mathbf{x}, \mathbf{e}_k) \mathbf{e}_k & \text{for } k = j, \\ 0 & \text{for } k \neq j. \end{cases} \end{aligned}$$

Note that it follows immediately from definition (5.70) that $\mathbf{P}_k$ commutes with every operator which commutes with $\mathbf{A}$.

From relations (5.68), (5.69) and (5.70) we get the following expressions for $\mathbf{x}$ and $\mathbf{Ax}$:

$$\mathbf{x} = \sum_{k=1}^n \mathbf{P}_k \mathbf{x}, \quad (5.71)$$

$$\mathbf{Ax} = \sum_{k=1}^n \lambda_k \mathbf{P}_k \mathbf{x}. \quad (5.71)$$

It follows from equation (5.71) that the operator $\sum_{k=1}^n \mathbf{P}_k$ is *identical*:

$$\mathbf{I} = \sum_{k=1}^n \mathbf{P}_k. \quad (5.73)$$

Equation (5.72) yields a so-called **expansion of a self-adjoint operator in terms of eigenfunctions**:

$$A = \sum_{k=1}^n \lambda_k P_k. \quad (5.74)$$

Properties 1° and 2° of projections and relation (5.74) yield the following expression for A^2 :

$$A^2 = \sum_{k=1}^n \lambda_k^2 P_k.$$

It is evident that in general

$$A^s = \sum_{k=1}^n \lambda_k^s P_k \quad (5.75)$$

for any integral positive s . A

Let us consider an arbitrary polynomial $p(\lambda) = \sum_{i=1}^n c_i \lambda^i$. By definition, it is

assumed that $p(A) = \sum_{k=1}^m c_k A^k$. Paying attention to relation (5.75), it is easy to get the following expression for $p(A)$:

$$p(A) = \sum_{i=1}^m p(\lambda_i) P_i. \quad (5.76)$$

Let us prove the following theorem.

Theorem 5.23. (Hamilton-Cayley Theorem). *If A is a self-adjoint operator and $p(\lambda) = \det(A - \lambda I)$ is a characteristic polynomial of that operator, then*

$$p(A) = 0.$$

Proof. Indeed, if A is a self-adjoint operator and λ_i are its eigenvalues, then, in accordance with Theorem 5.8, λ_i is a root of the characteristic equation, i.e. $p(\lambda_i) = 0$. From this and from (5.76) it follows that $p(A) = 0$. We have proved the theorem.

5.5.6. Positive operators. The m th-degree roots of an operator. The self-adjoint operator A is said to be positive if the relation

$$(Ax, x) \geq 0 \quad (5.77)$$

holds for any x from V .

If the operator A is positive and it follows from the condition $(Ax, x) = 0$ that $x = 0$, then A is said to be a **positive definite operator**.

Positive and positive definite operators are designated as $A \geq 0$ and $A > 0$ respectively.

Note the following simple statement.

Every eigenvalue of a positive (positive definite) operator is non-negative (positive).

This statement follows from simple arguments.

Let λ be an eigenvalue of the operator A . Then, in accordance with the lemma in 5.5.4, there is an element x , $\|x\| = 1$, such that $\lambda = (Ax, x)$.

From this and from (5.77) we find that $\lambda \geq 0$ for positive operators and $\lambda > 0$ for positive definite operators. We have proved the statement.

We introduce the concept of an **m th-degree root** (m being a natural number) **of an operator**.

Definition. A root of degree m of the operator A is an operator B such that $B^m = A$.

A root of degree m of the operator A is symbolized as $A^{1/m}$. It is natural to isolate a class of operators for which the operation of finding an m th-degree root would have sense. The following theorem provides a definite answer to this question.

Theorem 5.24. Let A be a positive self-adjoint operator $A \geq 0$. Then there is a positive self-adjoint operator $A^{1/m}$, $\tilde{A}^{1/m} \geq 0$, for any natural m .

Proof. We denote by λ_k the eigenvalues of the operator A , and let $\{e_k\}$ be an orthonormal basis of the eigenvectors. Next we denote by P_k the projection onto a one-dimensional subspace generated by e_k .

In accordance with what was said in the preceding subsection, we have an expansion (5.74) of the self-adjoint operator $\mathbf{A}$ in terms of eigenfunctions:

$$\mathbf{A} = \sum_{k=1}^m \lambda_k \mathbf{P}_k. \quad (5.74)$$

Since $\lambda_k = 0$ (see the statement we have just proved), we can introduce the following self-adjoint operator $\mathbf{B}$:

$$\mathbf{B} = \sum_{k=1}^n \lambda_k^{1/m} \mathbf{P}_k. \quad (5.78)$$

In accordance with (5.70), there holds a relation

$$(\mathbf{P}_k \mathbf{x}, \mathbf{x}) \geq 0$$

which yields the positivity of the operators $\mathbf{P}_k$ and that of the operator $\mathbf{B}$ (see (5.78)).

It follows from properties 1° and 2° of the projections $\mathbf{P}_k$ (see 5.5.5) that

$$\mathbf{B}^m = \sum_{k=1}^n \lambda_k \mathbf{P}_k. \text{ Comparing this expression for } \mathbf{B}^m \text{ with expression (5.74) for } \mathbf{A},$$

we find that $\mathbf{B}^m = \mathbf{A}$. We have found above that the operator $\mathbf{B}$ is positive. And this is what we wished to prove.

Remark 1. We note without proof that there exists a unique positive operator $\mathbf{A}^{1/m}$.

Remark 2. In the orthonormal basis $\{\mathbf{e}_k\}$ of the eigenvectors of the operator $\mathbf{A}$ the matrix of the operator $\mathbf{A}^{1/m}$ has the following form:

$$\begin{pmatrix} \lambda_1^{1/m} & 0 & \dots & 0 \\ 0 & \lambda_1^{1/m} & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & \lambda_n^{1/m} \end{pmatrix}$$

5.6. Reducing a Quadratic Form to the Sum of Squares

We shall discuss here the problem of choosing a basis in which a quadratic form (an invariant quadratic function of the coordinates of a vector; the exact definition will be given later on) has the simplest form.

Quadratic forms will be studied in Chapter 7. In particular, we shall consider there various methods of reducing quadratic forms to a sum of squares.

We introduce the concept of the so-called **Hermitian forms**.

Definition. The sesquilinear form $B(x, y)$ is called **Hermitian** if for any x and y there holds a relation

$$B(x, y) = \overline{B(y, x)}. \quad (5.79)$$

In accordance with the corollary of Theorem 5.11, any sesquilinear form $B(x, y)$ (a Hermitian form inclusive) can be uniquely represented as

$$B(x, y) = (Ax, y), \quad (5.80)$$

where A is a linear operator.

Let us prove the following two statements in which we shall elucidate the conditions under which a sesquilinear form is Hermitian.

Theorem 5.25. For the sesquilinear form $B(x, y)$ to be Hermitian, it is necessary and sufficient that the operator A in representation (5.80) of this form should be self-adjoint ($A=A^*$).

Proof. Indeed, if A is a self-adjoint operator, then, using the properties of a scalar product, we obtain

$$B(x, y) = (Ax, y) = (x, Ay) = \overline{(Ax, y)} = \overline{B(y, x)}.$$

Thus condition (5.79) is fulfilled, i.e. the form $B(x, y) = (Ax, y)$ is Hermitian.

And if the form $B(x, y) = (Ax, y)$ is Hermitian, we can again take the properties of a scalar product into account and get

$$(Ax, y) = B(x, y) = \overline{B(y, x)} = \overline{(Ay, x)} = (x, Ay).$$

Thus, we have $(Ax, y) = (x, Ay)$, i.e. the operator A is self-adjoint.

And this is what we set out to prove.

Theorem 5.26. *For the sesquilinear form $B(\mathbf{x}, \mathbf{y})$ to be Hermitian, it is necessary and sufficient that the function $B(\mathbf{x}, \mathbf{x})$ should be real.*

Proof. The form $B(\mathbf{x}, \mathbf{y})$ is Hermitian if and only if the linear operator A in representation (5.80) of this form is self-adjoint (see Theorem 5.25). But in accordance with Theorem 5.18, for the operator A to be self-adjoint, it is necessary and sufficient that the scalar product $(A\mathbf{x}, \mathbf{x})$ should be real for any $\mathbf{x}$. We have proved the theorem.

We shall now introduce the concept of a **quadratic form**.

Let $B(\mathbf{x}, \mathbf{y})$ be a Hermitian form.

The **quadratic form** corresponding to the form $B(\mathbf{x}, \mathbf{y})$ is a function $B(\mathbf{x}, \mathbf{x})$.

We shall now prove the following *theorem on the reduction of a quadratic form to a sum of squares (to a diagonal form)*.

Theorem 5.27. *Suppose $B(\mathbf{x}, \mathbf{y})$ is a Hermitian form defined on various vectors $\mathbf{x}$ and $\mathbf{y}$ of the n -dimensional Euclidean space V . Then this space contains an orthonormal basis $\{\mathbf{e}_k\}$ and there are real numbers λ_k such that for any $\mathbf{x}$ belonging to V the quadratic form $B(\mathbf{x}, \mathbf{x})$ can be represented as the following sum of the squares of the coordinates of the vector $\mathbf{x}$ in the basis $\{\mathbf{e}_k\}$:*

$$B(\mathbf{x}, \mathbf{x}) = \sum_{k=1}^n \lambda_k |\xi_k|^2. \quad (5.81)$$

Proof. The form $B(\mathbf{x}, \mathbf{y})$ being Hermitian, there is self-adjoint operator A , in accordance with Theorem 5.25, such that

$$B(\mathbf{x}, \mathbf{y}) = (A\mathbf{x}, \mathbf{y}). \quad (5.82)$$

Let us direct our attention to Theorem 5.21. By this theorem, we can indicate for the operator A an orthonormal basis $\{\mathbf{e}_k\}$ consisting of the eigenvectors of that operator. If λ_k are the eigenvalues of A and ξ_k are the coordinates of the vector $\mathbf{x}$ in the basis $\{\mathbf{e}_k\}$ so that

$$\mathbf{x} = \sum_{k=1}^n \xi_k \mathbf{e}_k, \quad (5.83)$$

then, using formula (5.12) and the relation $\mathbf{A}\mathbf{e}_k = \lambda_k \mathbf{e}_k^*$, we get the following expression for $\mathbf{A}\mathbf{x}$:

$$\mathbf{A}\mathbf{x} = \sum_{k=1}^n \lambda_k \xi_k \mathbf{e}_k. \quad (5.84)$$

From (5.83), (5.84) and the orthonormality of the basis $\{\mathbf{e}_k\}$ we obtain the following expression for $(\mathbf{A}\mathbf{x}, \mathbf{x})$:

$$(\mathbf{A}\mathbf{x}, \mathbf{x}) = \sum_{k=1}^n \lambda_k |\xi_k|^2.$$

(* This relation follows from the fact that λ_k and $\mathbf{e}_k$ are the eigenvalues and eigenvectors of the operator $\mathbf{A}$ respectively.)

From this expression and from relation (5.82) we get (5.81). We have thus proved the theorem:

Let us now prove an important theorem on simultaneous reduction of two quadratic forms to a sum of squares.

Theorem 5.28. *Suppose $A(\mathbf{x}, \mathbf{y})$ and $B(\mathbf{x}, \mathbf{y})$ are Hermitian forms defined on various vectors $\mathbf{x}$ and $\mathbf{y}$ of the n -dimensional linear space V . We assume that an inequality $B(\mathbf{x}, \mathbf{x}) > 0$ holds for all nonzero elements $\mathbf{x}$ belonging to V . Then there is a basis $\{\mathbf{e}_k\}$ in the space V such that the quadratic forms $A(\mathbf{x}, \mathbf{x})$ and $B(\mathbf{x}, \mathbf{x})$ can be represented as*

$$A(\mathbf{x}, \mathbf{x}) = \sum_{k=1}^n \lambda_k |\xi_k|^2, \quad (5.85)$$

$$B(\mathbf{x}, \mathbf{x}) = \sum_{k=1}^n |\xi_k|^2, \quad (5.86)$$

where λ_k are real numbers and fit are the coordinates of the vector $\mathbf{x}$ in the basis $\{\mathbf{e}_k\}$.

Proof. Since under an additional requirement that $B(\mathbf{x}, \mathbf{x}) > 0$ for $\mathbf{x} \neq 0$ the properties of a scalar product and those of the Hermitian form $B(\mathbf{x}, \mathbf{y})$ are similarly formulated, we can introduce a scalar product $(\mathbf{x}, \mathbf{y})$ of vectors in the linear space V setting

$$(\mathbf{x}, \mathbf{y}) = B(\mathbf{x}, \mathbf{y}). \quad (5.87)$$

Thus V is a Euclidean space with the scalar product (5.87).

By Theorem 5.27, we can indicate in V an orthonormal basis $\{\mathbf{e}_k\}$ and real numbers λ_k such that in this basis the quadratic form $A(\mathbf{x}, \mathbf{x})$ is represented as (5.85).

On the other hand, in any orthonormal basis the scalar product $(\mathbf{x}, \mathbf{x})$, equal to $B(\mathbf{x}, \mathbf{x})$ in accordance with (5.87), is equal to the sum of squares of the absolute values of the coordinates of the vector $\mathbf{x}$. Thus, the representation of $B(\mathbf{x}, \mathbf{x})$ in the form (5.86) is also well-founded. And this is what we had to prove.

5.7. Unitary and Normal Operators

We discuss here the properties of a significant class of operators acting in the Euclidean space V .

Definition 1. The linear operator U from $L(V, V)$ is called **unitary** if for any elements $\mathbf{x}$ and $\mathbf{y}$ belonging to V there holds a relation

$$(U\mathbf{x}, U\mathbf{y}) = (\mathbf{x}, \mathbf{y}). \quad (5.88)$$

In what follows, relation (5.88) is called the **condition under which an operator is unitary**.

Remark 1. It follows from condition (5.88) under which an operator is unitary that for any unitary operator U the equality $\|U\mathbf{x}\| = \|\mathbf{x}\|$ holds true.

Note the following *assertion*.

If λ is an eigenvalue of the unitary operator U , then $|\lambda| = 1$.

Indeed, if λ is an eigenvalue of U , then there is an element e such that $\|e\| = 1$ and $Ue = \lambda e$. From this and from Remark 1 it follows that $|\lambda| = \|\lambda e\| = \|Ue\| = \|e\| = 1$. We have proved the assertion.

Let us now prove the following theorem.

Theorem 5.29. *For the linear operator U , acting in the Euclidean space V , to be unitary, it is necessary and sufficient that the relation*

$$U^* = U^{-1} \quad (5.89)$$

hold true.

Proof. (1) Necessity. Let the operator U be unitary, i.e. condition (5.88) be fulfilled. Taking into account the definition of an adjoint operator U^* , we can rewrite this condition in the following form*.

$$(U^* Ux, y) = (x, y), \quad (5.90)$$

or, to put it otherwise, the equality

$$((U^* - I)x, y) = 0$$

is satisfied for any x and y .

* Recall that the operator U^* is said to be an adjoint of the operator U if for any z and y the relation $(z, Uy) = (U^*z, y)$ is satisfied. Setting $z = Ux$, we get (5.90).

Fixing any element x in this equation and considering y to be arbitrary, we find that the linear operator $U^*U - I$ acts according to the rule $(U^*U - I)x = 0$.

Consequently, $U^*U = I$. We can make sure by analogy that $UU^* = I$.

Thus U and U^* are operators inverse of each other, i.e. relation (5.89) is satisfied. We have proved the necessity of the theorem.

(2) Sufficiency. Suppose condition (5.89) is fulfilled. Then, evidently, $UU^* = U^*U = I$. Directing our attention to the definition of an adjoint operator and using the relation we have just written, we get, for any x and y , relations

$$(Ux, Uy) = (x, U^*Uy) = (x, Iy) = (x, y).$$

Thus condition (5.88) under which the operator is unitary is ful-

filled. Consequently, the operator U is unitary. We have thus proved the theorem. '

Remark 2. While proving the theorem, we have found that condition (5.88) under which the operator U is unitary and the condition

$$U^*U = UU^* = I \quad (5.91)$$

are equivalent. We can thus base the definition of a unitary operator on condition (5.91).

This condition can also be called the **condition under which the operator U is unitary**.

We shall now introduce the concept of a normal operator.

Definition 2. The linear operator A is called **normal** if the relation

$$A^*A = AA^* \quad (5.92)$$

holds true.

Considering condition (5.91) under which an operator is unitary and condition (5.92), we see that *any unitary operator is a normal operator*.

We shall need the following auxiliary assertion.

Lemma. Let A be a normal operator. Then the operator A and the operator A^* have a common element e such that $\|e\| = 1$ and the relations $Ae = \lambda e$ and $A^*e = \lambda e$ hold true.

Proof. Suppose λ is an eigenvalue of the operator A and $R_\lambda = \ker (A - \lambda I)$. In other words, R_λ is a set of all elements x such that $Ax - \lambda x = 0$.

Let us verify that if x belongs to R_λ then A^*x also belongs to R_λ . Indeed, if $Ax = \lambda x$ (i.e. $x \in R_\lambda$), then

$$A(A^*x) = A^*(Ax) = A^*(\lambda x) = \lambda(A^*x)$$

since A is a normal operator.

To put it otherwise, the vector A^*x is an eigenvector of the operator A and corresponds to the eigenvalue λ , i.e. belongs to R_λ .

Considering the operator A^* as an operator from R_λ into R_λ and using the inference from the corollary of Theorem 5.8 stating that every linear operator has an eigenvalue, we can assert that there is an element e in R_λ such that $\|e\| = 1$ and the relations $A^*e = \mu e$ and $Ae = \lambda e$ hold true.

Using these relations and the condition $\|e\| = 1$, we find $(Ae, e) = (\lambda e, e) = \lambda \|e\|^2 = \lambda$, $(e, A^*e) = (e, \mu e) = \bar{\mu} \lambda \|e\|^2 = \bar{\mu} \lambda$. Since $(Ae, e) = (e, A^*e)$, it evidently follows that $\lambda = \bar{\mu} \lambda$. We have proved the lemma.

Let us now prove the following theorem.

Theorem 5.30. *Let A be a normal operator. Then there is an orthonormal basis $\{e_k\}$ consisting of the eigenelements of the operators A and A^* .*

Proof. In accordance with the lemma we have just proved, the operators A and A^* have a common eigenelement e_1 , belonging to V , with $\|e_1\| = 1$. The eigenvalues of the operators A and A^* corresponding to e_1 are equal to λ_1 and $\bar{\lambda}_1$ respectively.

Suppose V_1 is an orthogonal complement of the element e_1 to complete the space V . In other words, V_1 is the collection of all x satisfying the condition $(x, e_1) = 0$.

We shall prove that if x belongs to V_1 , then Ax and A^*x belong to V_1 . Indeed, if $(x, e_1) = 0$, then

$$(Ax, e_1) = (x, A^*e_1) = (x, \bar{\lambda}_1 e_1) = \bar{\lambda}_1 (x, e_1) = 0,$$

i.e. $Ax \in V_1$. By analogy, if $(x, e_1) = 0$, then

$$(A^*x, e_1) = (x, Ae_1) = (x, \lambda_1 e_1) = \lambda_1 (x, e_1) = 0,$$

i.e. $A^*x \in V_1$.

Thus V_1 is an invariant subspace of the operators A and A^* . Therefore, by the lemma we have just proved, there is an eigenelement e_2 of the operators A and A^* in the subspace V_1 such that $\|e_2\|=1$, $Ae_2=\lambda_2e_2$, $A^*e_2=\bar{\lambda}_2e_2$.

Next we designate as V_2 an orthogonal complement of the element e_2 to complete V_1 . Reasoning as above, we shall prove that there is an eigenelement e_3 of the operators A and A^* in V_2 such that $\|e_3\|=1$. Continuing with similar arguments, we shall, evidently, construct, in the space V , an orthonormal basis $\{e_k\}$ consisting of the eigenelements of the operators A and A^* . We have proved the theorem.

Corollary 1. *Let A be a normal operator. There is a basis $\{e_k\}$ in which A has a diagonal matrix.*

Indeed, in accordance with the theorem we have proved, there is a basis $\{e_k\}$ consisting of the eigenvectors of the operator A . By Theorem 5.9, the matrix of the operator A is diagonal in this basis.

Corollary 2. *A unitary operator has a complete orthonormal system of eigenvectors.*

Here is the reciprocal of Theorem 5.30.

Theorem 5.31. *If the operator A acting in the n -dimensional Euclidean space V has n pairwise orthogonal eigenelements $e_1, e_2, \dots, e_n$, then the operator A is normal.*

Proof. Suppose $\{e_k\}$ are pairwise orthogonal eigenvectors of the operator A . Then $Ae_k=\lambda_k e_k$ and, according to Theorem (5.69), there holds the following representation of the operator A^* :

$$Ax = \sum_{k=1}^n \lambda_k (x, e_k) e_k.$$

We shall prove that the adjoint operator A^* acts according to the rule

$$A^* y = \sum_{k=1}^n \bar{\lambda}_k(y, e_k) e_k. \quad (5.93)$$

It is sufficient to prove that the equality

$$(x, A^* y) = (Ax, y) \quad (5.94)$$

is valid for the operators A and A^* defined by relations (5.69) and (5.93).

Substituting the expression $A^* y$ into the left-hand side of this equation in accordance with formula (5.93), we get, after simple transformations,

$$\begin{aligned} (x, A^* y) &= \sum_{k=1}^n (x, \bar{\lambda}_k(y, e_k) e_k) = \sum_{k=1}^n \lambda_k(y, e_k) \overline{(x, e_k)} \\ &= \sum_{k=1}^n \lambda_k(x, e_k) (e_k, y). \end{aligned}$$

We have thus proved equality (5.94) and, therefore, the operator A^* acting according to rule (5.93) is an adjoint of the operator A .

To complete the proof of the theorem, we must verify the validity of equality (5.92): $A^* A = A A^*$.

In accordance with (5.93)**, we have

$$\begin{aligned} A A^* x &= \sum_{k=1}^n \bar{\lambda}_k(x, e_k) A e_k \\ &= \sum_{k=1}^n \lambda_k \bar{\lambda}_k(x, e_k) e_k = \sum_{k=1}^n \bar{\lambda}_k \lambda_k(x, e_k) e_k = A^* A x. \end{aligned}$$

Thus, equation (5.92) holds for the operators A and A^* and, consequently, the operator A is normal. And this completes the proof of the theorem.

*Representation (5.69) is valid for any operator having n pairwise orthogonal eigenvectors.

** We have also used the relations $A e_k \lambda_k e_k$.

5.8. Canonical Form of Linear Operators

We consider here the choice of a special basis for a given linear operator, in which the matrix of the operator has the simplest form called the **Jordan form of a matrix**.

We introduce the concept of a **principal element** of the operator.

Definition. *The element x is called the **principal element** of the operator A corresponding to the eigenvalue λ if the relations*

$$(A - \lambda I)^m \neq 0, (A - \lambda I)^{m+1} x = 0$$

hold for a certain integral $m \geq 1$. The number m here is called the grade of the principal element x .

In other words, if x is a principal element of grade m , then the element $(A - \lambda I)^m x$ is an eigenvector of the operator A .

We shall prove here the following *fundamental* theorem.

Theorem 5.32. *Let A be a linear operator acting in the n -dimensional Euclidean space V . There is a basis*

$$\{ e_k^m \}, \quad k, = 1, 2, \dots, l, m = 1, 2, \dots, n_k;$$

$$n_1 + n_2 + \dots + n_l = n \quad (5.95)$$

which is formed from eigenvectors and principal vectors of the operator A in which the action of the operator A is described by the following relations:

$$A e_k^1 = \lambda_k e_k^1, \quad k, = 1, 2, \dots, l;$$

$$A e_k^m = \lambda_k e_k^m + e_k^{m-1}, \quad k = 1, 2, \dots, l; m = 1, 2, \dots, n_k. \quad (5.96)$$

Before proving the theorem, we shall make some remarks.

Remark 1; The vectors e_k^1 ($k = 1, 2, \dots, l$) evidently of the basis (5.95) are, evidently, the eigenvectors of the operator A which correspond to the eigenvalues λ_k .

From the definitions of principal vectors and relations (5.96) it follows that the vectors e_k^m ($k=1, 2, \dots, l; m=1, 2, \dots, n_k$) are principal vectors of grade m corresponding to the eigenvalues λ_k respectively.

Remark 2. Paying attention to formulas (5.13) and (5.12), we see that relations (5.96) indeed define the action of the operator A in the space V at the given basis $\{e_k^m\}$.

Remark 3. The matrix A of the linear operator A in the basis $\{e_k^m\}$ has the following “block” form:

$$A = \begin{pmatrix} \Lambda_1 & & & & 0 \\ & \Lambda_2 & & & \\ & & \ddots & & \\ & & & \ddots & \\ 0 & & & & \Lambda_l \end{pmatrix} \quad (5.97)$$

where the block Λ_k is the following matrix:

$$\Lambda_k = \begin{pmatrix} \lambda_k & 1 & 0 & \dots & 0 \\ 1 & \lambda_k & 1 & \dots & 0 \\ \dots & \dots & \dots & \dots & \dots \\ 0 & 0 & 0 & \dots & 1 \\ 0 & 0 & 0 & \dots & \lambda_k \end{pmatrix}. \quad (5.98)$$

Remark 4. Form (5.97) of the matrix A of the linear operator A is called the **Jordan form of the matrix** of that operator. The block Λ_k is usually called a **Jordan block** (or submatrix) of the matrix A . Note that Theorem 5.32 on the reduction of the matrix of an operator to the block form (5.97) is known as the theorem on the reduction of the matrix of an operator to the Jordan form.

Remark 5. The Jordan form of matrix (5.97) is defined with an accuracy to within the order of arrangement of the blocks Λ_k along the diagonal of the matrix. This order depends on the order of enumeration of the eigenvalues λ_k .

We present the proof of Theorem 5.32 suggested by A. F. Filippov. (A. F. Filippov. "TA Brief Proof of the Theorem on the Reduction of a Matrix to the Jordan Form". *Vestruk Moscovskogo Universiteta*, 1971, No. 2.)

Proof of Theorem 5.32. We shall prove the theorem by induction. For $n = 1$ the assertion of the theorem is obvious. Suppose $n > 1$ and the theorem is valid for spaces of dimension smaller than n . We shall prove that under this assumption it is also valid for spaces of dimension n . And this completes the proof of the theorem.

Let λ be an eigenvalue of the operator A . In accordance with Theorem 5.8, this number is a root of the characteristic equation $\det (A - \lambda I) = 0$. Consequently, the rank r of the linear operator

$$B = A - \lambda I \quad (5.99)$$

is smaller than n , i.e. $r < n$.

The linear operator B maps the space V onto the subspace $\text{im } B$. Therefore, the operator B maps the subspace $\text{im } B$ of dimension $r < n$ into the same subspace. By induction assumption, there is a basis

$$\{ h_k^m \}, k = 1, 2, \dots, p; m = 1, 2, \dots, r_k;$$

$$r_1 + r_2 + \dots + r_p = r, \quad (5.100)$$

in $\text{im } B$ in which the action of the operator B from $\text{im } B$ into $\text{im } B$ is given by the following relations :

$$\left. \begin{aligned} B h_k^1 &= \mu_k h_k^1, k = 1, 2, \dots, p, \\ B h_k^m &= \mu_k h_k^m + h_k^{m-1}, k = 1, 2, \dots, p; m = 2, 3, \dots, r_k. \end{aligned} \right\} \quad (5.101)$$

Thus, in this basis the matrix $\tilde{B}$ of the operator $\tilde{B}$ from $\text{im } B$ into $\text{im } B$ has the following block form:

$$\tilde{B} = \begin{pmatrix} M_1 & & & 0 \\ & M_2 & & \\ & & \ddots & \\ 0 & & & M_p \end{pmatrix}, \text{ where } M_k = \begin{pmatrix} \mu_k & 1 & 0 & \dots & 0 \\ 0 & \mu_k & 1 & \dots & 0 \\ \dots & \dots & \dots & \dots & \dots \\ 0 & 0 & 0 & \dots & 1 \\ 0 & 0 & 0 & \dots & \mu_k \end{pmatrix}. \quad (5.102)$$

Suppose that only the first m_1 ($m_1 \geq 0$) eigenvalues of the operator $\tilde{B}$ are zero.

Since the rank of each submatrix M_k (see (5.102)), for which $\mu_k = 0$, is equal to r_{k-1} and the rank of the submatrix, for which $\mu_k \neq 0$, is equal to r_k , we find, by

(5.100), that the rank of the matrix $\tilde{B}$ is equal to $\sum_{k=1}^p r_k - m_1 = r - m_1$. Therefore,

the dimension of the subspace $\ker \tilde{B}$ is a linear closure of the vectors $\mathbf{h}_1^1$ and $\mathbf{h}_2^1, \dots, \mathbf{h}_{m_1}^1$. By virtue of linear independence, these vectors form a basis in $\ker \tilde{B}$.

Evidently, $\ker \tilde{B} \subset \ker B$. Let us complete the basis $\mathbf{h}_1^1, \mathbf{h}_2^1, \dots, \mathbf{h}_{m_1}^1$ in $\ker \tilde{B}$ by the vectors $\mathbf{g}_k, k = 1, 2, \dots, m_0, m_0 = n - r - m_1$ to form a basis in $\ker B$ (by Theorem 5.1, the dimension of $\ker B$ is equal to $n - \dim \operatorname{im} B$, i.e. to $n - r$).

Since $\mathbf{g}_k \in \ker B$, we have

$$B\mathbf{g}_k = 0. \quad (5.103)$$

Let us now consider the vectors $\mathbf{h}_k^{r_k}, k = 1, 2, m_1$. since these vectors belong to $\operatorname{im} B$, there are vectors $\mathbf{f}_k \in V$ such that

$$B\mathbf{f}_k = \mathbf{h}_k^{r_k}, k = 1, 2, \dots, m_1. \quad (5.104)$$

We shall prove that the vectors

$$\left. \begin{array}{l} \mathbf{h}_k^m (k = 1, 2, \dots, p; \quad m = 1, 2, \dots, r_k, \\ \mathbf{g}_k (k = 1, 2, \dots, m_0; \quad \mathbf{f}_k (k = 1, 2, \dots, m_1) \end{array} \right\} \quad (5.105)$$

are linearly independent.

Let us consider the linear combination $\mathbf{f}$ of these vectors which is equal to zero:

$$\mathbf{f} = \sum_{k=1}^p \sum_{m=1}^{r_k} \alpha_{km} \mathbf{h}_k^m + \sum_{k=1}^{m_0} \beta_k \mathbf{g}_k + \sum_{k=1}^{m_1} \gamma_k \mathbf{f}_k = 0. \mathbf{e}_k \quad (5.106)$$

How does the operator $\mathbf{B}$ act on the element $\mathbf{f}$?

In accordance with (5.101), (5.103) and (5.104), we get the following expression:

$$\mathbf{B}\mathbf{f} = \sum_{k=1}^p \alpha_{k1} \mu_k \mathbf{h}_k^1 + \sum_{k=1}^p \sum_{m=2}^{r_k} \alpha_{km} (\mu_k \mathbf{h}_k^m + \mathbf{h}_k^{m-1}) + \sum_{k=1}^{m_1} \gamma_k \mathbf{h}_k^{rk} = 0. \quad (5.107)$$

Relation (5.107) is a linear combination of the base vectors $\{\mathbf{h}_k^m\}$ which is equal to zero; therefore, the coefficients in these vectors $\mu_k = 0$ in the indicated linear combination are zero. Since $\mu_k = 0$ for $k \geq m_1$, it follows from (5.107) that the coefficients in $\mathbf{h}_k^{rk}$ are exactly equal to γ_k and, therefore, $\gamma_k = 0$. From this and from relation (5.106) we get equalities

$$\mathbf{g} = \sum_{k=1}^{m_0} \beta_k \mathbf{g}_k = - \sum_{k=1}^p \sum_{m=1}^{r_k} \alpha_{km} \mathbf{h}_k^m, \quad (5.108)$$

from which it follows that the vector $\mathbf{g}$ which is a linear combination of vectors $\{\mathbf{g}_k\}$ belongs to $\ker \mathbf{B}$ (recall that the vectors $\{\mathbf{g}_k\}$ form a part of the basis in $\ker \mathbf{B}$).

On the other hand, it follows from (5.108) that $\mathbf{g}$ is a linear combination of vectors $\mathbf{h}_k^m$, i.e. belongs to $\text{im } \tilde{\mathbf{B}}$. Consequently, $\mathbf{g}$ belongs to $\ker \tilde{\mathbf{B}}$ (recall that $\ker \tilde{\mathbf{B}}$ is

an intersection of $\text{im } \tilde{\mathbf{B}}$ and $\ker \mathbf{B}$) and, therefore, $\mathbf{g} = \sum_{k=1}^{m_1} \delta_k \mathbf{h}_k^1$.

Since the linear closures of the collections of vectors $\{\mathbf{g}_k\}$ and $\{\mathbf{h}_k^1\}$ only have a zero element in common (taken together, these collections form a basis in $\ker \mathbf{B}$)

and, as we have established, $\mathbf{g}$ belongs to each of the indicated linear closures, we have $\mathbf{g} = 0$. But then it follows from (5.108) that $\beta_k = 0$ ($k = 1, 2, \dots, m_0$) and $\alpha_{km} = 0$ ($k = 1, 2, \dots, p; m = 1, 2, \dots, r_k$).

Thus, all the coefficients in the linear combination (5.106) of vectors (5.105) are equal to Zero, i.e. vectors (5.105) are linearly independent.

The total number of vectors (5.105) is equal to $r + m_0 + m_1$. Since $m_0 = n - r - m_1$ (we have established this fact above, while proving the theorem, when we introduced the vectors $\mathbf{g}_k$), the total number of vectors (5.105) is equal to n and, therefore, they form a basis in V . We introduce the designation

$$\mathbf{f}_k = \mathbf{h}_k^{r_k + 1} \quad (5.109)$$

and write the vectors forming this basis in the following sequence of series:

$$\left. \begin{aligned} &\{\mathbf{g}_1\}; \{\mathbf{g}_2\}; \dots; \{\mathbf{g}_{m_0}\}; \\ &\{\mathbf{h}_k^1, \dots, \mathbf{h}_k^{r_k}, \mathbf{h}_k^{r_k + 1}\}, k = 1, 2, \dots, m_1; \\ &\{\mathbf{h}_k^1, \dots, \mathbf{h}_k^{r_k}\}, k = m_1 + 1, \dots, p_1. \end{aligned} \right\} \quad (5.110)$$

Let us see how the operator $\mathbf{B}$ acts on the vectors of basis (5.110) in the space V . Considering relations (5.101), (5.103), (5.104) and (5.109), we make sure that the action of $\mathbf{B}$ in basis (5.110) is specified by the relations

$$\mathbf{B}\mathbf{g}_k = 0, k = 1, 2, \dots, m_0, \mathbf{B}\mathbf{h}_k^{r_k + 1} = \mathbf{h}_k^{r_k}, k = 1, 2, \dots, m_1.$$

and relations (5.101)._

Thus, in basis (5.110), the operator $\mathbf{B} = \mathbf{A} - \lambda \mathbf{I}$ acts according to rule (5.96), indicated in the formulation of Theorem 5.32. But then, in this basis, the operator $\mathbf{B} = \mathbf{A} + \lambda \mathbf{I}$ also acts by the same rule. We have proved the theorem.

5.9. Linear Operators in a Real Euclidean Space

We shall show here how the results and definitions presented in the previous sections can be extended to the case of real Euclidean spaces.

5.9.1. General remarks. Let us consider an arbitrary n -dimensional real Euclidean space V and an operator A from V into V .

For the case of a real linear space, the concept of a linear operator is formulated by a complete analogy with the respective concept for a complex space.

Definition 1. *The operator A is said to be **linear** if for any elements $x \in V$ and $y \in V$ and any real numbers α and β the following equality holds true:*

$$A(\alpha x + \beta y) = \alpha Ax + \beta Ay. \quad (5.111)$$

The concepts of an eigenvalue and an eigenvector of an operator are introduced by a complete analogy with a complex space.

It should be pointed out that eigenvalues are roots of the characteristic equation of an operator.

The converse statement in a real case is valid only when the corresponding root of the characteristic equation is real. Only in that case is the indicated root an eigenvalue of the linear operator in question.

In this connection, it is natural to isolate some class of linear operators in a real Euclidean space, all roots of whose characteristic equations are real.

We have proved in Theorem 5.16 that all the eigenvalues of a self-adjoint operator are real. In addition, the concept of a self-adjoint operator played an important part in the inferences concerning quadratic forms made in 5.6. It is, therefore, natural to extend the notion of a self-adjoint operator to the case of a real space.

We shall first introduce the concept of the operator A^* which is an adjoint of the operator A . Namely, the operator A^* is said to be an adjoint of A if for any x and y from V the equality $(Ax, y) = (x, A^*y)$ is satisfied.

Without any difficulty we can extend Theorem 5.12 on the existence and uniqueness of an adjoint operator to the case of a real space.

Recall that the proof of Theorem 5.12 is based on the concept of a sesquilinear form. In the real case, instead of a sesquilinear form use should be made of the bilinear form $B(x, y)$. A requisite remark was made in this connection in 5.4.2.

Let us recall, in this connection, the definition of a bilinear form in any real, not necessarily Euclidean, linear space L . Let B be a function putting every ordered pair $(\mathbf{x}, \mathbf{y})$ of vectors $\mathbf{x} \in L$ and $\mathbf{y} \in L$ into correspondence with the real number $B(\mathbf{x}, \mathbf{y})$.

Definition 2. *The function $B(\mathbf{x}, \mathbf{y})$ is called a bilinear form given on L if for any vectors $\mathbf{x}, \mathbf{y}$ and $\mathbf{z}$ from L and for any real number λ the relations*

$$\left. \begin{aligned} B(\mathbf{x} + \mathbf{z}, \mathbf{y}) &= B(\mathbf{x}, \mathbf{y}) + B(\mathbf{z}, \mathbf{y}), \\ B(\mathbf{x}, \mathbf{y} + \mathbf{z}) &= B(\mathbf{x}, \mathbf{y}) + B(\mathbf{x}, \mathbf{z}), \\ B(\lambda \mathbf{x}, \mathbf{y}) &= B(\mathbf{x}, \lambda \mathbf{y}) = \lambda B(\mathbf{x}, \mathbf{y}), \end{aligned} \right\} \quad (5.112)$$

hold true.

An important part will be played in this section by the special representation of the bilinear form $B(\mathbf{x}, \mathbf{y})$ as

$$B(\mathbf{x}, \mathbf{y}) = (A\mathbf{x}, \mathbf{y}), \quad (5.113)$$

where A is some linear operator. The corresponding theorem (Theorem 5.11) on an analogous representation of a sesquilinear form in a complex space was based on the inferences made in the lemma in 5.4.1 on a special representation of the linear form $f(\mathbf{x})$. It was indicated at the end of 5.4.1 that the lemma was true in a real space as well. It should only be pointed out that in the proof of the lemma the choice of the elements $\mathbf{h}^k$ must be performed not by formula (5.41) but by the formula $\mathbf{h}^k = f(\mathbf{e}_k)$, where $f(\mathbf{x})$ is a given linear form in a real space.

In 5.6 we introduced Hermitian forms. The Hermitian form is the sesquilinear form $B(\mathbf{x}, \mathbf{y})$ in a complex space characterized by the relation $B(\mathbf{x}, \mathbf{y}) = \overline{B(\mathbf{y}, \mathbf{x})}$ (the bar over B signifies that we take a complex conjugate value of B).

In the case of a real space, the analogs of Hermitian forms are symmetric bilinear forms. Such a form is characterized by the relation

$$B(\mathbf{x}, \mathbf{y}) = B(\mathbf{y}, \mathbf{x}). \quad (5.114)$$

The bilinear form $B(\mathbf{x}, \mathbf{y})$ given in the linear space L is said to be skew-symmetric if the relation $B(\mathbf{x}, \mathbf{y}) = -B(\mathbf{y}, \mathbf{x})$ holds true for any vectors $\mathbf{x}$ and $\mathbf{y}$ belonging to L . It is evident that for every bilinear form the functions

$$B_1(\mathbf{x}, \mathbf{y}) = \frac{1}{2} [B(\mathbf{x}, \mathbf{y}) + B(\mathbf{y}, \mathbf{x})],$$

$$B_2(\mathbf{x}, \mathbf{y}) = \frac{1}{2} [B(\mathbf{x}, \mathbf{y}) - B(\mathbf{y}, \mathbf{x})],$$

are a symmetric and a skew-symmetric form respectively. Since $B(\mathbf{x}, \mathbf{y}) = B_1(\mathbf{x}, \mathbf{y}) + B_2(\mathbf{x}, \mathbf{y})$, we obtain the following statement:

Any bilinear form can be represented as the sum of symmetric and skew-symmetric bilinear forms.

It is easy to see that such a representation is unique.

We shall prove the following theorem on symmetric bilinear forms (this theorem is an analog of Theorem 5.25 on Hermitian forms).

Theorem 5.33. *For the bilinear form $B(\mathbf{x}, \mathbf{y})$ specified on various vectors $\mathbf{x}$ and $\mathbf{y}$ of the real Euclidean space V to be symmetric, it is necessary and sufficient that the linear operator A appearing in representation (5.113) should be self-adjoint.*

Proof. If A is a self-adjoint operator, we can use the properties of a scalar product and get

$$B(\mathbf{x}, \mathbf{y}) = (A\mathbf{x}, \mathbf{y}) = (A\mathbf{y}, \mathbf{x}) = B(\mathbf{y}, \mathbf{x}).$$

Thus relation (5.114) holds true, i.e. the bilinear form $B(\mathbf{x}, \mathbf{y}) = (A\mathbf{x}, \mathbf{y})$ is symmetric.

Now if the form $B(\mathbf{x}, \mathbf{y}) = (A\mathbf{x}, \mathbf{y})$ is symmetric, then the relations $(A\mathbf{x}, \mathbf{y}) = B(\mathbf{x}, \mathbf{y}) = B(\mathbf{y}, \mathbf{x}) = (A\mathbf{y}, \mathbf{x})$ hold true.

Consequently, the operator A is self-adjoint. And this is what we wished to prove.

We introduce the concept of the matrix of the linear operator A . Let $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ be

some basis in the n -dimensional real linear space L . We set $A\mathbf{e}_k = \sum_{i=1}^n \alpha_k^i \mathbf{e}_i$.

Then, as in a complex case, it is easy to show that if $\mathbf{x} = \sum_{k=1}^n x^k \mathbf{e}_k$, then the

representation $y^i = \sum_{k=1}^n a_k^i x^k$ holds for the components of the vector $\mathbf{y} = A\mathbf{x}$.

The matrix $A(a_k^i)$ is known as the matrix of the linear operator A in the basis $\{e_k\}$.

By analogy with what was done in 5.2, we can prove that the quantity $\det A$ does not depend on the choice of the basis, and thus we correctly define $\det A$ of the operator A .

The *characteristic equation* corresponding to the operator A is an equation $\det(A - \lambda I) = 0$, and the polynomial appearing on the left-hand side of this equation is called a *characteristic polynomial of the operator A*.

Let us now prove the theorem on the roots of a characteristic polynomial. of a self-adjoint operator in a real Euclidean space.

Theorem 5.34. *All the roots of the characteristic polynomial of the self-adjoint linear operator A in a Euclidean space are real.*

Proof. Suppose $\lambda = \alpha + i\beta$ is a root of the characteristic equation

$$\det(A - \lambda I) = 0 \quad (5.115)$$

of the self-adjoint operator A .

We fix some basis $\{e_k\}$ in V and designate as a_{jk} the elements of the matrix of the operator A in that basis (note that a_{jk} are real numbers).

We shall seek a nonzero solution of the following system of homogeneous linear equations with respect to $\xi_1, \xi_2, \dots, \xi_n$:

$$\sum_{k=1}^n a_{jk} \xi_k = \lambda \xi_j, \quad j = 1, 2, \dots, n, \quad (5.116)$$

where $\lambda = \alpha + i\beta$.

Since the determinant of system (5.116) is equal to $\det(A - \lambda I)$ (recall that the determinant of the matrix of a linear transformation does not depend on the choice of the basis and, according to (5.115), that determinant is zero), system (5.116) of the homogeneous linear equations has a nonzero solution $\xi_k = x_k + iy_k$, $k = 1, 2, \dots, n$.

Substituting this solution into the right-hand and left-hand sides of system (5.116), taking into account that $\lambda = \alpha + i\beta$ and separating the real and the imaginary part of the relations obtained, we find that the collections $(x_1, x_2, \dots, x_n)$ and $(y_1, y_2, \dots, y_n)$ of real numbers* satisfy the following system of equations:

$$\left. \begin{aligned} \sum_{k=1}^n a_{jk} x_k &= \alpha x_j - \beta y_j, \\ \sum_{k=1}^n a_{jk} y_k &= \alpha y_j + \beta x_j, \quad j = 1, 2, \dots, n. \end{aligned} \right\} \quad (5.117)$$

Let us consider the vectors $\mathbf{x}$ and $\mathbf{y}$ with the coordinates $(x_1, x_2, \dots, x_n)$ and $(y_1, y_2, \dots, y_n)$, respectively, in the given basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$. Then we can rewrite relations (5.117) in the form

$$\mathbf{A}\mathbf{x} = \alpha\mathbf{x} - \beta\mathbf{y}, \quad \mathbf{A}\mathbf{y} = \alpha\mathbf{y} + \beta\mathbf{x}.$$

Multiplying scalarly the first of the relations obtained by $\mathbf{y}$ and the second relation by $\mathbf{x}$, we, evidently, obtain equalities

$$\left. \begin{aligned} (\mathbf{A}\mathbf{x}, \mathbf{y}) &= \alpha(\mathbf{x}, \mathbf{y}) - \beta(\mathbf{y}, \mathbf{y}), \\ (\mathbf{x}, \mathbf{A}\mathbf{y}) &= \alpha(\mathbf{x}, \mathbf{y}) + \beta(\mathbf{x}, \mathbf{x}). \end{aligned} \right\} \quad (5.118)$$

* Recall that not all these numbers are zero.

The operator $\mathbf{A}$ being self-adjoint, we have $(\mathbf{A}\mathbf{x}, \mathbf{y}) = (\mathbf{x}, \mathbf{A}\mathbf{y})$. Therefore, subtracting relations (5.118), we get an equation

$$\beta[(\mathbf{x}, \mathbf{x}) + (\mathbf{y}, \mathbf{y})] = 0.$$

But $(\mathbf{x}, \mathbf{x}) + (\mathbf{y}, \mathbf{y}) \neq 0$ (if $(\mathbf{x}, \mathbf{x}) + (\mathbf{y}, \mathbf{y}) = 0$, then $x_k = 0$ and $y_k = 0$, $k = 1, 2, \dots, n$; consequently, the solution $\xi_k = x_k + iy_k$ would be zero whereas, by construction, this solution is nonzero). Therefore, $\beta = 0$, and since β is an imaginary part of the root $\lambda = \alpha + i\beta$ of the characteristic equation (5.115), λ is, evidently, a real number. We have proved the theorem.

As in the complex case, the statement concerning the existence of an orthonormal basis consisting of the eigenvectors of that operator holds true for a self-adjoint operator (the analog of theorem 5.21). Let us prove this statement. 1

Theorem 5.35. *Every self-adjoint linear operator A acting in the n -dimensional real Euclidean space V possesses an orthonormal basis of eigenvectors.*

Proof. Let λ_1 be a real eigenvalue of the operator A , and e_1 is a unit eigenvector corresponding to that eigenvalue ($\|e_1\|=1$).

We designate as V_1 an $(n - 1)$ -dimensional subspace of the space V which is orthogonal to e_1 . V_1 is evidently an invariant subspace of the space V (i.e. if $x \in V_1$ then $Ax \in V_1$). Indeed, suppose $x \in V_1$; then $(x, e_1) = 0$. Since the operator A is self-adjoint and λ_1 is an eigenvalue of A , we get $(Ax, e_1) = (x, Ae_1) = \lambda_1 (x, e_1) = 0$.

Consequently, $Ax \in V_1$ and, therefore, V_1 is an invariant sub-space of the operator A . We can, therefore, consider the operator A in the subspace V_1 . It is clear that in V_1 the operator A is self-adjoint. By Theorem 5.34, the operator A acting in V_1 has a real eigenvalue λ_2 which is associated with an eigenvector $e_2 \in V_1$ of the operator A satisfying the condition $\|e_2\|=1$.

Considering then an $(n - 2)$ -dimensional subspace of V_2 which is orthogonal to the vectors e_1 and e_2 and repeating the arguments we have just presented, we construct an eigenvector e_3 of the operator A which is orthogonal to the vectors e_1 and e_2 and satisfies the condition $\|e_3\|=1$.

Continuing with the same reasoning, we find, as a result, n mutually orthogonal eigenvectors $e_1, e_2, \dots, e_n$ of the operator A satisfying the condition $\|e_k\|=1, k=1, 2, \dots, n$. The vectors $\{e_k\}$ evidently form a basis in V . And this is what we set out to prove.

Remark. Assume that $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ is an orthonormal basis in the n -dimensional Euclidean space V consisting of the eigenvectors of the self-adjoint operator A , i.e. $A\mathbf{e}_k = \lambda_k \mathbf{e}_k$. Then the matrix of the operator A in the basis $\{\mathbf{e}_k\}$ is diagonal and the diagonal elements have the form $a_k^k \lambda_k$.

Note that if $\{\mathbf{e}_k\}$ is an arbitrary orthonormal basis in the real Euclidean space V , then the matrix of the self-adjoint operator A is symmetric, i.e. $A' = A$. The converse statement is also true, i.e. if in a certain orthonormal basis $\{\mathbf{e}_k\}$ the matrix of an operator is symmetric, then the operator A is self-adjoint.

This is what distinguishes a real case from a complex one, since in a complex case the operator A is self-adjoint if and only if the matrix A of that operator is Hermitian in an orthonormal basis, i.e. the elements a_k^i of the matrix A satisfy the condition $a_k^i = \overline{a_i^{-k}}$ (the bar signifies complex conjugation).

The indicated assertion follows directly from the fact that if (a_k^i) is a matrix of the operator A , then the matrix of the conjugate operator is equal to (a_k^i) in a real case and to (a_k^{-k}) in a complex case, which can be easily verified by direct calculations.

5.9.2. Orthogonal operators. In a complex Euclidean space an important role is played by unitary operators introduced in 5.7. Orthogonal operators are the analog of unitary operators in a real Euclidean space.

Definition 1. The linear operator P acting in the real Euclidean space V is said to be **orthogonal** if for any $\mathbf{x}$ and $\mathbf{y}$ belonging to V the equality

$$(P\mathbf{x}, P\mathbf{y}) = (\mathbf{x}, \mathbf{y}) \quad (5.119)$$

holds true.

Thus, an orthogonal operator retains a scalar product. Hence it follows directly that if $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ is an orthonormal basis of the Euclidean space V , then $P\mathbf{e}_1, P\mathbf{e}_2, \dots, P\mathbf{e}_n$ is also an orthonormal basis. In what follows, we shall call condition (5.119) the condition of orthogonality of the operator P .

The following statement holds true.

Theorem 5.36. *For the linear operator P to be orthogonal, it is necessary and sufficient that there exist an operator P^{-1} and there hold an equation*

$$P^* = P^{-1}, \quad (5.120)$$

where P^* is an adjoint operator of P , and P^{-1} is an operator inverse with respect to P .

Proof. (1) Necessity. Let P be an orthogonal operator, i.e. condition (5.119) be satisfied. Applying the adjoint operator P^* , we can rewrite this condition as $(P^*Px, y) = (x, y)$.

Thus, for any x and y , the equality $((P^*P - I)x, y) = 0$ holds true.

Fixing any element x in this equality and considering y to be arbitrary, we find that the linear operator $P^*P - I$ acts according to the rule $(P^*P - I)x = 0$.

Consequently, $P^*P = I$; we can verify by a complete analogy that $PP^* = I$. Thus the operators P^* and P are inverse of each other, i.e. condition (5.120) is fulfilled.

(2) Sufficiency. Assume that condition (5.120) is fulfilled. Then, evidently, $PP^* = P^*P = I$.

Directing our attention to an adjoint operator and using the relations we have just written, we get, for any x and y relations

$$(Px, Py) = (x, P^*y) = (x, Iy) = (x, y).$$

We see that condition (5.119) is fulfilled and, consequently, the operator P is orthogonal. We have thus proved the theorem.

We shall introduce now the concept of an orthogonal matrix P .

Definition 2. *The matrix P is said to be orthogonal if*

$$P'P = PP' = I, \quad (5.121)$$

where P' is a transposed matrix and I is an identity matrix.

If $e_1, e_2, \dots, e_n$ is an orthonormal basis in the Euclidean space V , then the operator P is orthogonal if and only if its matrix in the basis $\{e_k\}$ is orthogonal.

It follows directly from equation (5.121) that if the matrix $P = (p_i^k)$ is orthogonal then

$$\sum_{i=1}^n p_i^k p_i^l = \begin{cases} 1 & \text{for } k = l, \\ 0 & \text{for } k \neq l. \end{cases}$$

In a complex Euclidean space a unitary matrix is an analog of an orthogonal matrix. Namely, the matrix U is said to be *unitary* if there holds a relation

$$U^*U = UU^* = I, \quad (5.122)$$

in which U^* is a Hermitian conjugate matrix, i.e. $U = \bar{U}'$, where the prime means a transposition and the bar, a complex conjugation.

It is easy to show that in an orthonormal basis the matrix of the linear operator U is unitary if and only if the operator U is unitary.

In conclusion we shall consider an example of an orthogonal transformation in one-dimensional and two-dimensional spaces.

In a one-dimensional case, every vector $\mathbf{x}$ has the form $\mathbf{x} = \alpha \mathbf{e}$, where α is a real number and $\mathbf{e}$ is a vector generating the given space. Then $\mathbf{P}\mathbf{e} = \lambda \mathbf{e}$, and since $(\mathbf{P}\mathbf{e}, \mathbf{P}\mathbf{e}) = \lambda^2 (\mathbf{e}, \mathbf{e}) = (\mathbf{e}, \mathbf{e})$, we have $\lambda = \pm 1$.

Thus, in a one-dimensional case there are two orthogonal transformations, $\mathbf{P} + \mathbf{x} = \mathbf{x}$ and $\mathbf{P} - \mathbf{x} = -\mathbf{x}$.

In a two-dimensional case, every orthogonal transformation is defined in an arbitrary orthonormal basis by an orthogonal matrix of order 2, i.e. the matrix

$P = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$. From the condition $\mathbf{P}\mathbf{P}' = \mathbf{P}'\mathbf{P} = I$ it follows that

$$a^2 + b^2 = 1, a^2 = d^2, b^2 = c^2, ac + db = 0, ab + cd = 0.$$

Setting $a = \cos \phi$, $b = -\sin \phi$, we find that every orthogonal matrix order 2 has the form

$$P_{\pm} = \begin{pmatrix} \cos \phi & -\sin \phi \\ \pm \sin \phi & \pm \cos \phi \end{pmatrix}.$$

In both cases, we must take either the plus sign or the minus sign in the second row. Note that $\det P_{\pm} = \pm 1$. The orthogonal matrix P_{+} is said to be proper and the orthogonal matrix P_{-} , *improper*.

The operator P_{+} with the matrix P_{+} performs, in the orthonormal basis e_1, e_2 , a rotation in the plane e_1, e_2 , through the angle ϕ .

To find out how the operator P_{-} with the matrix P_{-} acts, we introduce the matrix $Q = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$ coinciding with P_{-} for $\phi = 0$, and note that $P_{-} = QP_{+}$. The matrix Q is associated with the reflection of the plane with respect to the e_1 axis and, consequently, the action of the operator P_{-} consists in a rotation through the angle $(\phi$ and in subsequent reflection.

Note that the vectors $P_{+}e_1, P_{+}e_2$ form an orthonormal basis by virtue of the orthogonality of P_{+} and in that basis the matrix of the operator P_{-} coincides with Q , i.e. is diagonal.

In a general case, when the orthogonal operator P acts in an n -dimensional Euclidean space, there is an orthonormal basis $e_1, e_2, \dots, e_n$ in which the matrix of the operator P has the form

$$\begin{pmatrix} 1 & & & & & & & & & & \\ & 1 & & & & & & & & & \\ & & \ddots & & & & & & & & \\ & & & \ddots & & & & & & & \\ & & & & -1 & & & & 0 & & \\ & & & & & -1 & & & & & \\ & & & & & & \cos \phi_1 - \sin \phi_1 & & & & \\ & & & & & & \sin \phi_1 - \cos \phi_1 & & & & \\ & & & & & & & \ddots & & & \\ & & & & & & & & \ddots & & \\ & 0 & & & & & & & & \ddots & \\ & & & & & & & & & & \cos \phi_k - \sin \phi_k \\ & & & & & & & & & & \sin \phi_k - \cos \phi_k \end{pmatrix}$$

In this matrix, all the elements, except for those written out, are equal to zero.

Thus, in a certain orthonormal basis, the action of an orthogonal operator reduces to successive rotations and reflections with respect to the coordinate axes.

Chapter – 6

ITERATIVE METHODS OF SOLVING LINEAR SYSTEMS AND EIGENVALUE PROBLEMS

In this chapter we study various methods of solving systems of linear equations with real coefficients with respect to the unknowns which also assume real values.

All methods of solving systems of linear equations which find practical applications can be divided into two large groups, **exact** methods and **iterative** methods.

An **exact** method is a method making it theoretically possible to obtain exact values of the unknowns as a result of a **finite** number of arithmetical operations. Cramer's rule, presented in Chapter 3, can serve as an example of an exact method.

Iterative methods enable us to find a solution only as a limit of a sequence of vectors which are constructed by means of a uniquely defined process called a process of iterations (successive approximations). Iterative methods are very convenient for the use of modern computers.

The first section of the present chapter is devoted to the exposition of the most applicable iterative methods of solving linear systems.

Iterative methods also find wide application in solving another very important computing problem of linear algebra, the so-called **complete problem of eigenvalues** (this is the name used for the problem of seeking all eigenvalues and the corresponding eigenvectors of a given matrix). In iterative methods the eigenvalues are calculated as the limits of some numerical sequences without first determining the coefficients in a characteristic polynomial.

In 6.2 we shall discuss one of the most important methods of solving a complete problem of eigenvalues used in computing technique, the so-called **rotation method** (or Jacobi's method for computing eigenvalues). This method can be applied to any symmetric (or Hermitian) matrix, can be easily realized on computers and always converges. It is stable as far as the errors of rounding-off the results of intermediary calculations are concerned and possesses a remarkable property that

the existence of multiple eigenvalues, close to one another, not only does not retard its convergence but, on the contrary, accelerates it. Jacobi's method known from the middle of the nineteenth century did not find practical application for a long time because of a large number of calculations necessary for its realization. Only the advent of fast electronic computers turned it into the most efficient method of solving the complete problem of eigenvalues of symmetric and Hermitian matrices.

6.1. Iterative Methods of Solving Linear Systems

6.1.1. The method of simple iteration (Jacobi's method). Let us consider a quadratic system of linear equations with real coefficients (3.10) (see 3.2.1), which we shall write in matrix form

$$AX = F, \quad (6.1)$$

assuming A to be the matrix of the coefficients of the linear equations

$$A = \begin{vmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \cdot & \cdot & \cdot & \cdot \\ a_{n1} & a_{n2} & \cdots & a_{nn} \end{vmatrix} \quad (6.2)$$

and X and F to be the vector-columns of the form

$$X = \begin{vmatrix} x_1 \\ x_2 \\ \cdot \\ x_n \end{vmatrix}, \quad F = \begin{vmatrix} f_1 \\ f_2 \\ \cdot \\ f_m \end{vmatrix}$$

the first of which must be defined and the second one is given.

Assuming a single-valued solvability of system (6.1), we replace the matrix equation (6.1) by the equivalent matrix equation $X = X - \tau AX + \tau F$ in which τ denotes a real number known as a **stationary parameter**.

With the aid of the last equation, we form an iterative sequence of the vectors $\{ X_k \}$ defining it by the recurrence relation

$$X_{k+1} = X_k - \tau AX_k + \tau F \quad (k = 0, 1, \dots) \quad (6.3)$$

with an arbitrary choice of the “zero” approximation X_0 .

The method of simple iteration consists in replacing an exact solution X of system (6.1) by the k th iteration X_k with a sufficiently large number k . Let us estimate the **error** $Z_k = X_k - X$ of the method of simple iteration.

Relations (6.3) and (6.1) immediately yield the following matrix equation for the error Z_k :

$$X_{k+1} = (I - \tau A)Z_k, \quad (6.4)$$

where I is an identity matrix of order n .

Let us consider the norm of a vector in the space E^n and the operator form of a square matrix of order n . As is customary, we use the term the **norm of the vector** X for the number $\|X\|$ which is equal to the square root of the sum of the squares of the coordinates of that vector. The **operator norm** of the arbitrary **matrix** A is the number $\|A\|$ equal to either the least upper bound of the relation $\|AX\| / \|X\|$ on the set of all nonzero vectors X or (which is the same thing) the greatest lower bound of the norms $\|AX\|$ on the set of all vectors X having a norm equal to unity.

Thus, by definition,

$$\|A\| = \sup_{X \neq 0} \frac{\|AX\|}{\|X\|}. \quad (6.5)$$

Recall that for any symmetric matrix A^* (* The matrix A is said to be symmetric if $A = A'$.) the operator norm of that matrix is equal to the eigenvalue of this matrix with the greatest absolute value (see 5.5.4), i.e.

$$\|A\| = \sup_s |\lambda_s|. \quad (6.6)$$

Equation (6.5) yields the following inequality valid for any matrix A and any vector X :

$$\|AX\| \leq \|A\| \|X\|. \quad (6.7)$$

From the matrix equation for an error (6.4) and from inequality (6.7) we find that for any number k

$$\|Z_{k+1}\| \leq \|I - \tau A\| \|XZ_k\|. \quad (6.8)$$

Let us now prove the following simple but important theorem.

Theorem 6.1. *For the iterative sequence (6.3) to converge to the exact solution X of system (6.1) for any choice of the zero approximation X_0 and for the given value of the parameter τ , it is sufficient that the condition*

$$\rho = \|I - \tau A\| < 1 \quad (6.9)$$

be satisfied. In this case, sequence (6.3) converges with the speed of a geometric progression with the common ratio ρ .

When the matrix A is symmetric, condition (6.9) is also a necessary condition for the convergence of the iterative sequence (6.3) for any choice of the zero approximation X_0 .

Proof. To establish the sufficiency of condition (6.9), we note that inequality (6.8) yields the following relation:

$$\|Z_k\| \leq \|I - \tau A\|^k \|X_0\|. \quad (6.10)$$

It is evident from (6.10) that condition (6.9) ensures the convergence of the sequence of errors Z_k to zero with the speed of a geometric progression with the common ratio ρ .

When the matrix A is symmetric, the matrix $I - \tau A$ is also symmetric and, therefore, by virtue of (6.6), condition (6.9) can be rewritten in the equivalent form

$$\rho = \sup_s |1 - \tau \lambda_s| < 1 \quad (6.11)$$

(the symbol $\{\lambda_s\}$ denotes here the eigenvalues of the matrix A).

Let us verify the fact that condition (6.11) is the necessary condition for the convergence of the sequence $\{Z_k\}$ to zero for any choice of the zero

approximation X_0 . We assume that condition (6.11) is not satisfied. Then there is an eigenvalue λ_s satisfying inequality $|1 - \tau\lambda_s| \geq 1$. We denote by $X^{(s)}$ the eigenvector of the matrix A corresponding to that eigenvalue and choose a zero approximation X_0 such that Z_0 coincides with $X^{(s)}$. Then; writing successively relation (6.4) for the numbers $1, 2, \dots, k$, we find that $Z_k = (1 - \tau\lambda_s)^k Z_0$. By virtue of the inequality $|1 - \tau\lambda_s| \geq 1$, it follows from the last relation that $\|Z_k\|$ does not tend to zero as $k \rightarrow \infty$. We have proved Theorem 6.1.

We draw attention to the fact that for practical applications it is not enough to establish the fact of the convergence of the sequence of iterations. The central problem of numerical methods is the estimation of the speed of convergence. It is very important to know the best way to use the stationary parameter τ to obtain the fastest convergence. Let us discuss this problem in more detail.

Suppose we are given an ε -accuracy with which we must obtain an exact solution of system (6.1). It is required to find an iteration X_k with a number k such that

$$\|Z_k\| \leq \varepsilon \|Z_0\|. \quad (6.12)$$

It follows from (6.9) and (6.10) that $\|Z_k\| \leq \rho^k \|Z_0\|$ and, therefore, (6.12) holds true for $\rho^k \leq \varepsilon$, i.e. for $k \geq \frac{\ln(1/\varepsilon)}{\ln(1/\rho)}$.

This shows that to decrease the number of iterations k sufficient to achieve the required ε -accuracy, we must choose the parameter τ such that the minimum of the function $\rho = \rho(\tau) = \|I - \tau A\|$ is attained.

Assuming the matrix A to be *symmetric* and positive definite, we arrive at the following optimization problem: it is required to find the minimum of the function.

$$\min_{\tau} \rho(\tau) = \|I - \tau A\| = \min_{\tau} \{ \max_s |1 - \tau\lambda_s| \}$$

The solution of this and of a somewhat more general problem suggested by A.A. Samarsky is presented in 6.1.2. We shall prove there that the indicated minimum of

the function $\rho = \rho(\tau)$ can be attained for the value $\tau = \frac{2}{(\gamma_1 + \gamma_2)}$, where γ_1 and γ_2 are the minimum and the maximum eigenvalues of the matrix A respectively, the minimum value of the function $\rho(\tau)$ being equal to

$$\frac{1 - \frac{\gamma_1}{\gamma_2}}{1 + \frac{\gamma_1}{\gamma_2}} = \frac{\gamma_2 - \gamma_1}{\gamma_2 + \gamma_1}.$$

6.1.2. The general implicit method of simple iteration. We shall again consider the solution of linear system (6.1), but this time we replace the iterative sequence (6.3) by a more general iterative sequence defined by the relation

$$B \frac{X_{k+1} - X_k}{\tau} + AX_k = F, \quad (6.13)$$

in which B is some “easily invertible” square matrix of order n and τ is a stationary parameter. This method of forming an iterative sequence is called an **implicit method of simple iteration**. The explicit method of simple iteration considered in 6.1.1 can be obtained from the implicit method in the special case $B = I$, where I is an identity matrix of order n .

To formulate the condition of convergence of the general implicit method of simple iteration in a form convenient for application, we must recall certain notions introduced in the previous chapter.

Recall that the matrix A is called *positive definite* if $(AX, X) > 0$ for any nonzero vector X . It was proved in Chapter 5 that the necessary and sufficient condition for the positive definiteness of the symmetric matrix A (or, what is the same, of a self-adjoint linear operator A) was the positiveness of all the eigenvalues of that matrix (that operator).

If the matrix A is positive definite, then we agree to write the inequality $A > 0$. Next we agree to write the inequality $B > A$ (or $A < B$) when $B - A > 0$ (i.e. when the matrix $B - A$ is positive definite).

We shall prove the following remarkable theorem*.(** This theorem is a special case of a considerably more general statement proved by the eminent Soviet mathematician A. A. Samarsky.*)

Theorem 6.2 (Samarsky's theorem). *Let the matrix A be symmetric and the conditions $A > 0$, $B > 0$ be fulfilled (symmetry of the matrix B is, in general, not assumed).*

Then, for the iterative sequence defined by the relations

$$B \frac{X_{k+1} - X_k}{\tau} + AX_k + F \quad (6.13)$$

to converge to the exact solution X of the system $AX = F$ for any choice of the zero approximation X_0 , it is sufficient that the following conditions should be satisfied:

$$2B > \tau A, \tau A > 0. \quad (6.14)$$

Under an additional assumption that the matrix B is symmetric, conditions (6.14) are not only sufficient but also necessary for the convergence of the indicated iterative-sequence for any choice of the zero approximation X_0 .

Proof. (1) Sufficiency. Let us first estimate the error $Z_k = X_k - X$. Since X satisfies the equation $AX = F$ and X_k satisfies relation (6.13), we get, for Z_k , a relation

$$B \frac{Z_{k+1} - Z_k}{\tau} + AZ_k = 0. \quad (6.15)$$

Let us establish for the error Z_k the so-called basic **energy relation**.

Multiplying (6.15) scalarly by the vector

$$2(Z_{k+1} - Z_k) = 2\tau \frac{Z_{k+1} - Z_k}{\tau},$$

we obtain an equation

$$2\tau = \left(\frac{Z_{k+1} - Z_k}{\tau}, \frac{Z_{k+1} - Z_k}{\tau} \right) + 2\tau \left(AZ_k, \frac{Z_{k+1} - Z_k}{\tau} \right) = 0. \quad (6.16)$$

It we use the designation $C = 2B - \tau A$ and the relation

$$Z_k = \frac{Z_{k+1} + Z_k}{2} - \frac{Z_{k+1} - Z_k}{2} = \frac{Z_{k+1} - Z_k}{2} - \frac{\tau}{2} \cdot \frac{Z_{k+1} - Z_k}{\tau},$$

we can rewrite equation (6.16) in the form

$$\tau \left(C \frac{Z_{k+1} - Z_k}{\tau}, \frac{Z_{k+1} - Z_k}{\tau} \right) + (A(Z_{k+1} + Z_k), Z_{k+1} - Z_k) = 0. \quad (6.11)$$

Next, we note that the matrix A being symmetric, the second term in (6.17) is equal to $(AZ_{k+1}, Z_{k+1}) - (AZ_k, Z_k)$. This leads us to the basic **energy relation**

$$\tau \left(C \frac{Z_{k+1} - Z_k}{\tau}, \frac{Z_{k+1} - Z_k}{\tau} \right) + (AZ_{k+1}, Z_{k+1}) = (AZ_k, Z_k). \quad (6.18)$$

To prove the sufficiency of condition (6.14), we must use the basic energy relation to prove the convergence of the sequence $\{\|Z_k\|\}$ to zero.

It follows from the basic energy relation and from the positive definiteness of the matrix $C = 2B - \tau A$ that $(AZ_{k+1}, Z_k) \leq (AZ_k, Z_k)$, i.e. it follows that the sequence $\{\|AZ_k, Z_k\|\}$ is non-increasing. In addition, it follows from the condition $A > 0$ that this sequence is bounded below by a zero and is, therefore, convergent..

But then it follows from the basic energy relation that

$$\lim_{k \rightarrow \infty} \left(C \frac{Z_{k+1} - Z_k}{\tau}, \frac{Z_{k+1} - Z_k}{\tau} \right) = 0. \quad (5.19)$$

Recall that for the positive definite matrix C there is always a $\delta > 0$ such that $(CX, X) \geq \delta (X, X)$ for any vector X or, what is the same, $\|X\|^2 \leq \frac{1}{\delta} (CX, X)$.

The last inequality enables us to infer that from the equality of limit (6.19) to zero it follows that

$$\lim_{k \rightarrow \infty} \|Z_{k+1} - Z_k\| = 0. \quad (6.20)$$

To complete the proof of the sufficiency, we must use the relation

$$B \frac{Z_{k+1} - Z_k}{\tau} + AZ_k = 0$$

from which, by virtue of the existence of a limited inverse matrix A^{-1} of the positive definite matrix A , it follows that

$$Z_k = -A^{-1} \cdot \frac{B}{\tau} (Z_{k+1} - Z_k).$$

The last equation and relation (6.20) make it legitimate to conclude that $\lim_{k \rightarrow \infty} \|Z_k\| = 0$. We have proved the sufficiency.

To prove the necessity of conditions (6.14) under an additional assumption that the matrix B is symmetric, we use the following lemma.

Lemma. *Let C be some symmetric matrix and B , a symmetric positive definite matrix. Then the matrix C is positive definite if and only if all the eigenvalues of the problem $CX = \lambda BX$ are positive.*

To prove the lemma, we note that since the matrix B is symmetric and positive definite, there is a self-adjoint positive definite operator $B^{1/2}$, by virtue of Theorem 5.24 from 5.5.6, such that the relation $B^{1/2} \cdot B^{1/2} = B$ holds for the corresponding matrix $B^{1/2}$. Since the matrix $B^{1/2}$ is symmetric and positive definite, there is an inverse bounded and symmetric matrix which we denote by $B^{-1/2}$.

Next we note that by means of the substitution $X = B^{-1/2} \cdot Y$ and the pre multiplication by the matrix $B^{-1/2}$, the problem on the eigenvalues $CX = \lambda BX$ passes into an *equivalent* problem on the eigenvalues $B^{-1/2} \cdot C \cdot B^{1/2} Y = \lambda Y$, so that to prove the lemma, it is sufficient to ascertain that the obviously symmetric matrix

$B^{-1/2} \cdot C \cdot B^{-1/2}$ is positive definite, if and only if the matrix C is positive definite. This latter follows immediately from the fact that for any two nonzero vectors X and Y , related as $Y = B^{-1/2} \cdot X$, the equality

$$(B^{-1/2} \cdot C \cdot B^{-1/2} X, x) = (C \cdot B^{-1/2} X, B^{-1/2} X) (CY, Y)$$

holds true. We have proved the lemma. .

We can now pass to the proof of the necessity of conditions (6.14) of Theorem 6.2 under an additional assumption that the matrix B is symmetric.

(2) Necessity. We proceed from the following assertion from the lemma proved above: *if the matrix B is symmetric and positive definite and the matrix C is symmetric and not positive definite, then the problem on the eigenvalues $CX = \lambda BX$ has at least one non- positive eigenvalue λ_s .*

Assume that the first of the conditions (6.14) is not satisfied, i.e. the requirement $2B - \tau A > 0$ is not met.

Setting $C = 2B - \tau A$ in the arguments presented above, we find

that the problem on eigenvalues $(2B - \tau A) X = \lambda BX$ has at least one nonpositive eigenvalue λ_s . We designate as $X^{(s)}$ the eigenvector corresponding to λ_s and choose a zero approximation x_0 such that the condition $Z_0 = X^{(s)}$ is fulfilled.

Then, rewriting the equation for the error (6.15) as $BZ_{k+1} = -BZ_k + (2B - \tau A) Z_k$, we get, setting successively, $k = 0, 1, \dots, Z_1 = (-1 + \lambda_s) X^{(s)}, Z_2 = (-1 + \lambda_s)^2 X^{(s)}, \dots, Z_k = (-1 + \lambda_s)^k X^{(s)}, \dots$. Since $-1 + \lambda_s \leq -1$, it evidently follows that $\{Z_k\}$ does not tend to zero as $k \rightarrow \infty$.

The case when the second of conditions (6.14) is not satisfied, i.e. the condition $\tau A > 0$ is not fulfilled, can be treated by analogy. In that case, in the arguments presented above we must set $C = \tau A$. We shall find that the problem $\tau AX = \lambda BX$ has at least one non-positive eigenvalue λ_s with the eigenvector $X^{(s)}$. Choosing a

zero approximation Z_0 such that the equality $Z_0 = X^{(s)}$ holds true and rewriting (6.15) in the equivalent form $BZ_{k+1} = BZ_k - \tau AZ_k$, we find that

$$Z_1 = (1 - \lambda_s) X^{(s)}, Z_2 = (1 - \lambda_s)^2 X^{(s)}, \dots, Z_k = (1 - \lambda_s)^k X^{(s)}, \dots$$

Since $\lambda_s \leq 0$, it is evident that $\{Z_k\}$ does not tend to zero as $k \rightarrow \infty$. And this completes the proof of Theorem 6.2.

Let us now estimate the speed of convergence of the general implicit method of simple iteration. Proceeding in the same way as Samarsky, we decide on the choice of the value of the parameter τ which would ensure the fastest convergence.

Assume that the matrix B is symmetric and positive definite. With the aid of such a matrix it is natural to introduce the so-called “**energy**” **scalar product** of two arbitrary vectors X and Y setting it equal to $(BX, Y) = (X, BY)$. We shall symbolize such a scalar product as $(X, Y)_B$.

With the aid of the matrix $B^{1/2}$ we can write this scalar product as $(X, Y)_B = (B^{1/2}, B^{1/2}X, Y) = (B^{1/2}X, B^{1/2}Y)$. Using the last equation, we can easily verify the validity of the four axioms of a scalar product (see 4.1.1) for the scalar product we have introduced.

Next it is natural to introduce the **energy norm** of the vector X setting it equal to $\sqrt{(X, Y)_B} = \sqrt{(BX, X)}$. We designate this energy norm as $\|X\|_B$.

Two distinct norms $\|X\|_I$ and $\|X\|_II$ of the same collection of vectors are said to be **equivalent** if there are positive constants λ_1 and λ_2 such that the inequalities

$$\lambda_1 \|X\|_I \leq \|X\|_{II} \leq \lambda_2 \|X\|_I$$

hold true.

It should be noted that the *energy norm of the vector X and its ordinary norm are equivalent*. Indeed, the validity of the inequality $\lambda_1 \|X\| \leq \|X\|_B$, i.e. the inequality $\lambda_1^2 (X, X) \leq (BX, X)$ follows from positive definiteness of the matrix

B , and the validity of the inequality $\|X\|_B \lambda_2 \|X\|$, i.e. the inequality $(BX, X) \leq \lambda_2^2 \|X\|^2$ follows from the Cauchy-Bunyakovsky inequality and from estimate (6.7) (it is sufficient to set $\lambda_2^2 = \|B\|$).

The equivalence of the ordinary and energy norms we have established makes it possible to assert that the *sequence $\|X_k\|$ converges to zero if and only if the sequence $\|X_k\|_B$ converges to zero.*

For further reasoning, the energy norm is more convenient than the ordinary norm. Let us prove the following fundamental theorem.

Theorem 6.3 (Samarsky's theorem). *Suppose the matrices A and B are symmetric and positive definite, and Z_k denotes the error of the general implicit method of simple iteration. Then, for the inequality $\|Z_k\|_B \leq \rho^k \|Z_0\|_B$, to hold for $\rho < 1$, it is sufficient that the condition*

$$\frac{1-\rho}{\tau} B \leq A \leq \frac{1+\rho}{\tau} B \quad (6.21)$$

holds true.

Remark. Samarsky proved that condition (6.21) was not only sufficient but also necessary for the validity of the inequality $\|Z_k\|_B \leq \rho^k \|Z_0\|_B$, but we shall not dwell on this fact here.

Proof of Theorem 6.3. For the sake of convenience, we divide the proof into two steps.

1°. We shall first prove that if the symmetric and positive definite matrices A and B satisfy Samarsky's conditions (6.14), then $(BZ_{k+1}, Z_{k+1}) \leq (BZ_k, Z_k)$.

Multiplying scalarly equation (6.15) by $2\tau Z_{k+1} = \tau(Z_{k+1} + Z_k) + \tau(Z_{k+1} - Z_k)$, we obtain

$$(B(Z_{k+1} - Z_k), Z_{k+1} + Z_k) + (B(Z_{k+1} - Z_k), Z_{k+1} - Z_k) + \tau(AZ_{k+1},$$

$$Z_{k+1} + Z_k + \tau (AZ_k, Z_{k+1} - Z_k) = 0.$$

In the last equation we replace AZ_k by the difference

$$\frac{1}{2} A(Z_{k+1} + Z_k) - \frac{1}{2} (Z_{k+1} - Z_k).$$

Then, taking into account the equality $(A(Z_{k+1} - Z_k), Z_{k+1} + Z_k) = (Z_{k+1} - Z_k, A(Z_{k+1} + Z_k))$ following from the symmetry of the matrix A , we get an identity

$$\begin{aligned} (B(Z_{k+1} - Z_k), Z_{k+1} + Z_k) + ((B - \frac{\tau}{2} A)(Z_{k+1} - Z_k), Z_{k+1} + Z_k) \\ + \frac{1}{2} (\tau A(Z_{k+1} + Z_k), Z_{k+1} + Z_k) = 0. \end{aligned}$$

Taking into consideration that (by virtue of Samarsky's condition (6.14)) the operators τA and $\tau B - \frac{\tau}{2} A$ are positive definite, we get from the last identity, an inequality

$$(B(Z_{k+1} - Z_k), Z_{k+1} + Z_k) \leq 0.$$

This inequality is equivalent to the inequality $(BZ_{k+1}, Z_{k+1}) \leq (BZ_k, Z_k)$ proved (by virtue of the identity $(BZ_{k+1}, Z_k) = (Z_{k+1}, BZ_k)$) following from the symmetry of the operator B .

2°. Suppose now that Samarsky's conditions (6.21) hold true for $\rho < 1$. We shall prove the validity of the inequality $\|Z_k\|_B \leq \rho^k \|Z_0\|_B$.

We set $Z_k = \rho^k V_k$. Then, evidently,

$$Z_{k+1} - Z_k = \rho^{k+1} V_{k+1} - \rho^k V_k = \rho^{k+1} (V_{k+1} - V_k) - (1 - \rho) \rho^k V_k.$$

Substituting these values of Z_k and $Z_{k+1} - Z_k$ into (6.15) and cancelling by ρ^k , we get the following relation for V_k :

$$\tilde{B} \frac{V_{k+1} - V_k}{\tau} + \tilde{A} V_k = 0. \quad (6.22)$$

in which $\tilde{B} = \rho B$, $\tilde{A} = A - \frac{1-\rho}{\tau} B$.

By virtue of conditions (6.21), the operators $\tilde{B}$ and $\tilde{A}$ satisfy the conditions $\tau \tilde{A} > 0$, $2\tilde{B} > \tau \tilde{A}$. From these conditions and from the fact that equation (6.22) for V_k is equivalent to (6.15) for Z_k , the first step for V_k yields the following estimate: $(\tilde{B} V_{k+1}, V_{k+1}) \leq (\tilde{B} V_k, V_k)$. Taking into account that $B = \rho B$, we get, from this estimate, an inequality $(B V_{k+1}, V_{k+1}) \leq (B V_k, V_k)$.

A successive application of the indicated inequality for the numbers $k = 0, 1, \dots$ leads us to the relation $(B V_k, V_k) \leq (B V_0, V_0)$, and the multiplication of the last relation by ρ^{2k} leads to the final estimate* $(B Z_k, Z_k) \leq \rho^{2k} (B V_0, V_0)$.

We have thus proved the inequality $\|Z_k\|_B \leq \rho^k \|Z_0\|_B$. And this completes the proof of Theorem 6.3.

In conclusion, we shall apply Samarsky's theorem 6.3 to decide on the choice of the value of the parameter τ for which the speed of the convergence is maximum. From the estimate $\|Z_k\|_B \leq \rho^k \|Z_0\|_B$ which was proved in Theorem 6.3 it follows that this problem reduces to finding the value of ρ for which the minimum value of the function $\rho = \rho(\tau)$ is attained.

Since both matrices A and B are symmetric and positive definite, there are positive constants λ_1 and λ_2 such that the inequalities $\lambda_1 B \leq A \leq \lambda_2 B$ hold true. We assume that the constants λ_1 and λ_2 in these inequalities are given**. Comparing the inequalities we have just written with conditions (6.21), we find that the minimum value of ρ is attained under the condition $\frac{(1-\rho)}{\tau} = \lambda_1$, $\frac{(1+\rho)}{\tau} = \lambda_2$,

whence we get the optimal value $\tau = \frac{2}{(\gamma_1 - \gamma_2)}$ and the minimum value of ρ which

is equal to $\frac{(\gamma_2 - \gamma_1)}{(\gamma_2 + \gamma_1)}$.

The explicit method of simple iteration studied in 6.1.1 is a special case of the consideration presented here. All the results we have obtained are valid for that method.

In the following three subsections, we shall consider, with the aid of the implicit method of simple iteration and Samarsky's theorem 6.2, some most applicable iteration methods and establish the conditions for the convergence.

6.1.3. A modified method of simple iteration. This method can be obtained from the general implicit method of simple iteration in the special case when the stationary parameter τ is equal to unity and the matrix B is a diagonal matrix D consisting of the elements of the matrix A lying on the principal diagonal, i.e. $B = D$, where

$$D = \begin{vmatrix} a_{11} & 0 & \dots & 0 \\ 0 & a_{22} & \dots & 0 \\ \cdot & \cdot & \cdot & \cdot \\ 0 & 0 & \dots & a_{nn} \end{vmatrix} \quad (6.23)$$

In this case, we assume, of course, that the matrix A is symmetric and that all its diagonal elements $a_{11}, a_{22}, \dots, a_{nn}$ are positive (the last requirement is necessary and sufficient for the diagonal matrix $B = D$ to be positive definite).

* We keep in mind that $z_k - \rho^k V_k, Z_0 = V_0$.

** It is natural to call the constants λ_1 and λ_2 the **equivalence constants** of the matrices A and B . For the commuting matrices A and B the constants λ_1 and λ_2 are equal, respectively, to the least and the greatest eigenvalue of the problem $AX = \lambda BX$.

It immediately follows from Theorem 6.2 that for the convergence of the modified method of simple iteration, for any choice of a zero approximation, it is sufficient that the conditions $2D > A$ and $A > 0$ be satisfied.

Theorem 6.1 makes it possible to express the sufficient condition for the convergence of the modified method of simple iteration in another form:

$$\|I - D^{-1}A\| < 1^* \quad (6.24)$$

(as before, the norm of the matrix is understood to be an operator form).

(* We take into account that in the case being considered we must take the matrix $\tilde{A}$ defined by the formula $\tilde{A} = B^{-1}A$ instead of the matrix A and set $B = D$, $\tau = 1$.)

Since $\|I - D^{-1}A\| = \|D^{-1}(D - A)\| = \|D^{-1}(A - D)\|$, the sufficient condition for convergence (6.24) can be rewritten in the equivalent form

$$\|D^{-1}(D - A)\| < 1. \quad (6.25)$$

Inequality (6.25) permits us to obtain various sufficient conditions for the convergence of the modified method of simple iteration.

Note, first of all, that if alongside the operator norm of matrix (6.2.) (which we denote by $\|A\|$ as before), we introduce the so-called spherical norm of that matrix which is designated as $\|A\|_{\text{sph}}$ and defined by the equation

$$\|A\|_{\text{sph}} = \left[\sum_{i=1}^n \sum_{j=1}^n a_{ij}^2 \right]^{1/2}, \text{ then, as was proved in 4.4 (see formula (4.28)), the}$$

inequality

$$\|AX\| \leq \|A\|_{\text{sph}} \|A\| \quad (6.26)$$

is valid for any vector X of the space E^n .

It immediately follows from (6.26) and (6.5) that the operator and the spherical norm of the matrix are related as $\|A\| \leq \|A\|_{\text{sph}}$.

Thus, by virtue of (6.25), the sufficient condition for the convergence of the modified method of simple iteration can be expressed by the inequality

$$\|D^{-1}(a-d)\|_{\text{sph}} < 1 \text{ which can be written out as } \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \frac{a_{ij}^2}{a_{ii}^2} < 1.$$

It should next be pointed out that when defining the operator norm (6.5) of the matrix A we proceeded from the ordinary (so-called **spherical**) norm of the vector

$$X = \left\| \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \right\| \text{ which is equal to } \|X\| = \left[\sum_{i=1}^n x_i^2 \right]^{1/2}.$$

Another Two norms of the vector X are often introduced: the so-called **cubic norm** defined by the equation $\max_{1 \leq i \leq n} |x_i|$ and the so-called **octahedral norm** defined by

the equation $\|X\|_{\text{oct}} = \sum_{i=1}^n |x_i|$. If the norm of the vector in definition (6.5) of the

operator norm of the matrix A is understood as its cubic or octahedral norm respectively, then relation (6.5) leads to the definition of the **cubic** and the **octahedral operator norm** of the matrix A respectively.

We can prove that the cubic and the octahedral operator norm of matrix (6.2) can be expressed in terms of the elements of that matrix as follows :

$$\|X\|_{\text{cub}} = \max_{1 \leq i \leq n} \sum_{j=1}^n |a_{ij}|,$$

$$\|X\|_{\text{oct}} = \max_{1 \leq j \leq n} \sum_{i=1}^n |a_{ij}|.$$

The verbatim repetition of the reasoning for the replacement of spherical norms by cubic and octahedral norms respectively leads to the sufficient condition for the convergence of the modified method of simple iteration expressed by relation (6.25) in which the norm of the matrix should be understood as its cubic or octahedral operator norm respectively.

This leads to the following two conditions each of which is a sufficient condition for the convergence of the modified method of simple iteration:

$$\sum_{\substack{j=1 \\ j \neq i}}^n \left| \frac{a_{ij}}{a_{ii}} \right| < 1 \quad (\text{for } i = 1, 2, \dots, n),$$

$$\sum_{\substack{i=1 \\ i \neq j}}^n \left| \frac{a_{ij}}{a_{ii}} \right| < 1 \quad (\text{for } j = 1, 2, \dots, n).$$

6.1.4. Seidel's method. Let us represent the symmetric matrix (6.2) as the sum of three matrices $A = D + L + U$, where D is the diagonal matrix (6.23) and L and U are, respectively, the strictly left and strictly right matrices having the form

$$L = \begin{bmatrix} 0 & 0 & \dots & 0 \\ a_{21} & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & 0 \end{bmatrix}, \quad U = \begin{bmatrix} 0 & a_{12} & \dots & a_{1n} \\ 0 & 0 & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{bmatrix}$$

and satisfying the condition $L' = U$.

Seidel's method can be obtained from the general implicit method of simple iteration in the special case when the stationary parameter τ is equal to unity and the matrix B is equal to the sum $D + L$. Thus the successive iterations in Seidel's method are defined by the relation

$$(D + L)(X_{k+1} - X_k) + AX_k = F.$$

Let us prove that *Seidel's method converges for any symmetric and positive definite matrix A .*

By Theorem 6.2, it is sufficient to prove that the condition

$$2(D + L) > A \tag{6.27}$$

holds for any such matrix A .

To prove (6.27), we pay attention to the fact that

$$(2(D + L)X, X) = (DX, X) + (DX, X) + (LX, X) + (LX, X)$$

$$\begin{aligned}
&= (DX, X) + (DX, X) + (LX, X) + (X, UX) \\
&= (DX, X) + (LX, X)
\end{aligned}$$

for any vector X .

Thus, to prove inequality (6.27), it is sufficient to verify that the matrix D is positive definite but this follows immediately from the fact that all the elements of the symmetric and positive definite matrix A which lie on the principal diagonal are positive. We have thus proved the convergence of Seidel's method.

6.1.5. The method of overrelaxation. This method can be derived from the general method of simple iteration in the special case when $\tau = \omega$, $B = D + \omega L$ and the parameter ω is chosen such that the eigenvalue of the matrix $I - \omega, (D + \omega L)^{-1} A$, carrying out the passage from the k th iteration to the $(k + 1)$ th, with the largest absolute value, is the least.

Let us prove that *if the matrix A is symmetric and positive definite, then for the overrelaxation method to converge, it is sufficient that the condition $0 < \omega < 2$ be satisfied.*

By Theorem 6.2, for the method to converge, it is sufficient that the conditions $\omega > 0$, $2(D + \omega L) > \omega A$ be satisfied.

For any vector X , the second condition leads to the inequality

$$(2(D + \omega L)X, X) > (\omega AX, X). \quad (6.28)$$

The last inequality is equivalent to each of the inequalities in the following chain:

$$\begin{aligned}
&(2DX, X) + (\omega LX, X) + (\omega LX, X) > (\omega AX, X), \\
&(2 - \omega)(DX, X) + (\omega DX, X) + (\omega LX, X) + (X, \omega UX) > (\omega AX, X), \\
&(2 - \omega)(DX, X) > 0.
\end{aligned}$$

We infer from the last inequality and from the positive definiteness of D that (6.28) is valid for $2 - \omega > 0$, i.e. for $\omega < 2$. We have thus proved that the conditions $0 < \omega < 2$ ensure the convergence of the method of overrelaxation.

6.1.6. The case of a nonsymmetric matrix A . When the matrix A is nonsymmetric, we can pre multiply matrix equation (6.1) by the matrix $\tilde{A}$ and replace equation (6.1) by the equation $\tilde{A}X = \tilde{F}$ in which $\tilde{F} = A'F$, $\tilde{A} = A'A$ so that the matrix $\tilde{A}$ will be symmetric and (as can be easily verified) positive definite.

6.1.7. Chebyshev's iterative method. Everywhere above, when we

considered the general implicit method of simple iteration, we supposed that the iteration parameter τ assumed one and the same constant value. An idea naturally arises to consider a more general case when the values of the iteration parameter in the indicated method depend on the number k of iteration. In that case, the sequence of the iterations is defined not by the relation

$$B \frac{X_{k+1} - X_k}{\tau} + AX_k = F, \quad (6.13)$$

but by a more general relation

$$B \frac{X_{k+1} - X_k}{\tau_{k+1}} + AX_k = F.$$

As before, B is a certain easily convertible square matrix of order n . With such a choice of the iteration sequence, we get the following relation for the error $Z_k = X_k - X$ of the iteration scheme:

$$B \frac{X_{k+1} - X_k}{\tau_{k+1}} + AZ_k = 0. \quad (6.15^*)$$

Let us assume that both matrices A and B are symmetric and positive definite. Then, as was noted above, there are positive constants γ_1 and γ_2 such that $\gamma_1 B \leq A \leq \gamma_2 B$. We assume that these constants are specified and recall once again that they are equal to the least and the greatest eigenvalue of the problem $AX = \lambda BX$ respectively. Let us evaluate the energy norm of the error $\|Z_k\|_B$.

Recall once again that for the symmetric and positive definite matrix B there is a symmetric and positive definite matrix $B^{1/2}$ such that $B^{1/2} \square B^{1/2} = B$. As before, we agree to use the symbol $B^{1/2}$ for the inverse of the matrix $B^{1/2}$.

To estimate the norm of the error Z_k , we make a replacement setting $B^{1/2} \square Z_k \cdot V_k$. With such a replacement, the relation for the error Z_k passes into the following relation for Z_k :

$$V_{k+1} = (I - \tau_{k+1} C) \cdot V_k \quad (k = 0, 1, 2, \dots),$$

where C denotes a matrix of the form $C = B^{-1/2} \square A \square B^{-1/2}$. Let us verify that the square of the ordinary norm of the vector V_k is equal to the square of the energy norm of the vector z_k . Indeed,

$$\|V_k\|^2 = (V_k, V_k) = (B^{1/2} Z_k, B^{1/2} Z_k) = (B Z_k, Z_k) = \|Z_k\|_B^2.$$

Thus, to estimate the energy norm Z_k , it is sufficient to estimate the ordinary norm V_k .

Let us estimate the norm $\|V_k\|$. We note, first of all that by means of the replacement $X = B^{-1/2} \square Y$, we can get inequalities $\gamma_1(Y, Y) \leq (CY, Y) \leq \gamma_2(Y, Y)$ from the inequalities $\gamma_1(BX, X) \leq (AX, X) \leq \gamma_2(BX, X)$. The inequalities $\gamma_1(Y, Y) \leq (CY, Y) \leq \gamma_2(Y, Y)$ are equivalent to $\gamma_1 I \leq C \leq \gamma_2 I$. Since, in addition, the matrix $C = B^{-1/2} \square B^{-1/2}$ is symmetric, all the eigenvalues of that matrix are real and lie on the interval $[\gamma_1, \gamma_2]$. Writing consecutively the relation $V_{k+1} (I - \tau_{k+1} \square C) V_k$ for the numbers $k = 0, 1, \dots$, we arrive at an equation

$$V_k = \prod_{j=1}^k (I - \tau_j C) \square V_0$$

which immediately yields

$$\|V_k\| \leq \left\| \prod_{j=1}^k (I - \tau_j C) \right\| \|V_0\|.$$

But then it follows from equation $\|V_k\| = \|Z_k\|_B$ that $\|Z_k\|_B \leq q^k = \|Z_0\|_B$, where $q_k = \left\| \prod_{j=1}^k (I - \tau_j C) \right\|$. Consequently, the iterative process converges provided that the sequence $\{q_k\}$ tends to zero, and the convergence is the faster the smaller the quantities q_k .

Since every value of q_k is a function of the parameters $\tau_1, \tau_2, \dots, \tau_k$, a problem arises of constructing an optimal collection of iteration parameters from the condition of minimum of q_k for a fixed k . Let us pass to the solution of this problem.

We assume that all the eigenvalues λ_s of the matrix C lie on the given interval $[\gamma_1, \gamma_2]$. Taking the symmetry of the matrix C into account, we arrive at the following optimization problem : it is required to find

$$\begin{aligned} \min_{(\tau_j)} q_k(\tau_1, \tau_2, \dots, \tau_k) &= \min_{(\tau_j)} \left\| \prod_{j=1}^k (I - \tau_j C) \right\| \\ &= \min_{(\tau_j)} \left\{ \max_s \left| \prod_{j=1}^k (I - \tau_j \lambda_s) \right| \right\}. \end{aligned}$$

Since all λ_s lie on the interval $[\gamma_1, \gamma_2]$, by extending the domain over which the maximum is taken, we find that

$$\min_{(\tau_j)} q_k(\tau_1, \tau_2, \dots, \tau_k) \leq \min_{(\tau_j)} \left\{ \max_{\lambda_1 \leq t \leq \lambda_2} \left| \prod_{j=1}^k (I - \tau_j t) \right| \right\}.$$

The rough problem we have obtained has a simpler solution. Besides, when solving such a problem, we do not use the information on the actual position of the

eigenvalues λ_s the interval $[\gamma_1, \gamma_2]$ but only take into consideration the boundaries of that interval. Such an approach makes it possible to construct a collection of optimal parameters for matrices of arbitrary structure.

Let us solve the rough problem of optimization. We set $P(t) = \prod_{j=1}^k (1 - \tau_j t)$ and note that the polynomial $P(t)$ satisfies the condition of valuation $P(0) = 1$. By changing the variable

$$t = \frac{1}{2} [\lambda_2 + \lambda_1 - S(\gamma_2 - \gamma_1)] = \frac{\lambda_2 + \lambda_1}{2} \left(1 - S \frac{\lambda_2 - \lambda_1}{\lambda_2 + \lambda_1} \right) = \frac{1 - S\rho_0}{\tau_0},$$

where $\rho_0 = \frac{\lambda_2 - \lambda_1}{\lambda_2 + \lambda_1}$, $\tau_0 = \frac{2}{\lambda_2 + \lambda_1}$, we map the interval $\gamma_1 \leq t \leq \gamma_2$ into the interval $-1 \leq S \leq 1$, the point $t = 0$ passing into the point $S = S_0 = \frac{1}{\rho_0} > 1$.

With such a change, the optimization problem in question passes into the following problem: *among all the polynomials $\bar{P}_k(S)$ of degree k which satisfy the valuation*

condition $\bar{P}_k\left(\frac{1}{\rho_0}\right) = 1$, find the polynomial for which $\max_{|S| \leq 1} |\bar{P}(S)|$ is minimal.

As we know, this is Chebyshev's polynomial $P_k(S) = T_k(S) \cdot [T_k(S_0)]^{-1}$, where

$$T_k(S) = \begin{cases} \cos(k \arccos S) & \text{for } |S| \leq 1, \\ \frac{1}{2} [(S + \sqrt{S^2 - 1})^k + (S - \sqrt{S^2 - 1})^k] & \text{for } |S| > 1. \end{cases}$$

Since $\max_{|S| \leq 1} |T_k(S)| = 1$, we have

$$\min_{(\tau_j)} \max_{\gamma_1 \leq t \leq \gamma_2} |P_k(t)| = \frac{1}{T_k(S_0)}, k \geq 1$$

with $\frac{1}{T_k(S_0)} = q_k \frac{2\rho_1^k}{1 + \rho_1^{2k}}$, where $\rho_1 = \frac{\sqrt{\gamma_2} - \sqrt{\gamma_1}}{\sqrt{\gamma_2} + \sqrt{\gamma_1}}$.

To calculate the optimal collection of parameters, we proceed from the equation

$$P_k(t) = \prod_{j=1}^k (1 - \varepsilon_j t) = q_k T_k \left(\frac{1 - \varepsilon_0 t}{\rho_0} \right)$$

(we have taken into account that $S = \frac{1 - \tau_0 t}{\rho_0}$). We equate the roots of the polynomials appearing on the left-hand and right-hand sides of this equation. Since the polynomial $P_k(t)$ has the roots $t_j = 1/\tau_j$ ($j=1, 2, \dots, k$) and the polynomial $T_k(S)$ has the roots $S_j = \cos \frac{2j-1}{2k} \pi$ ($j=1, 2, \dots, k$), we take into account that $t = \frac{1 - \rho_0 S}{\tau_0}$ and get $\frac{1}{\tau_j} = \frac{1 - S_j \rho_0}{\tau_0}$ ($j=1, 2, \dots, k$; S_j were defined above).

Thus, the values $\tau_j = \frac{\tau_0}{1 + S_0 S_j}$, where $S_j = \cos \frac{2j-1}{2k} \pi$, $j=1, 2, \dots, k$ are the optimal values of the iteration parameters.

The iteration process with the indicated collection of parameters is known as **Chebyshev's process**.

We arrive at the following theorem.

Theorem 6.4. *If the matrices A and B are symmetric and positive definite and if $\gamma_1 B \leq A \leq \gamma_2 B$, then Chebyshev's iterative process converges, and the estimate*

$$\|Z_k\|_B \leq q_k \|Z_0\|_B,$$

where $q_k = \frac{2\rho_1^k}{1 + \rho_1^{2k}}$ for $\rho_1 = \frac{\sqrt{\gamma_2} - \sqrt{\gamma_1}}{\sqrt{\gamma_2} + \sqrt{\gamma_1}}$ is valid for the error Z_k after k iterations.

If we take the requirement $\|Z_k\|_B < \varepsilon \|Z_0\|_B$ as the condition for ending the process for the preassigned ε -accuracy, then Theorem 6.4 yields the following

estimate for k iterations: $k \geq k_0(\epsilon) = \frac{\ln(\epsilon/2)}{\ln \rho_1}$. Comparing this estimate with the

obtained estimate $\bar{k} \geq \bar{k}_0(\epsilon) = \frac{\ln \epsilon}{\ln \rho_0}$ of the number of iterations for the method of

simple iteration, we find, provided that the quantity $\xi = \frac{\gamma_2}{\gamma_1}$ is small, that

$k_0(\epsilon) \approx \frac{\ln(2/\epsilon)}{2\sqrt{\xi}}$, $\bar{k}_0(\epsilon) \approx \frac{\ln(1/\epsilon)}{2\xi}$. The comparison of these two estimates shows

that Chebyshev's method has certain advantages (when the quantity $\xi = \frac{\gamma_1}{\gamma_2}$ is small).

Chebyshev's method we have described have been known since the beginning of 1950s. It is sometimes called **Richardson's method**.

It should be pointed out that we have studied the method for an ideal computing process with an infinite number of digits whereas calculations on computers are carried out with a finite number of digits and, therefore, there are numbers which are computer infinity M_∞ and computer zero. Ii, in the process of computer calculations, there appears a number M exceeding M_∞ an emergency shut-down of the computer occurs.

From the point of view of an ideal computing process, the values of the iteration parameters τ_j can be ordered in any way (by any of the $k!$ methods). From the viewpoint of an ideal computing process, any two sequences of iteration parameters $\{\tau_j\}$ are equivalent since the ϵ -accuracy they require can be attained in the same number of iterations.

But when calculations are carried out with the aid of computers, various sequences of parameters $\{\tau_j\}$ are not equivalent. For some sequences of values $\{\tau_j\}$ an emergency shut-down of the computer may occur because of a growth of intermediate values. For other sequences of values $\{\tau_j\}$ no emergency shut-down

occurs, but because of the nonmonotonic tendency of the error Z_k to zero, i.e. owing to the fact that the norm of the matrix $I - \tau_j C$ of transition from the $(j - 1)$ th iteration to the j th iteration can be greater than unity, the estimate established for an ideal situation is not valid for this error.

Because of the indicated circumstances, a theoretical problem arises which consists in indicating the best law of ordering the values $\{ \tau_j \}$ for which the influence of the errors of rounding-off is the least for Chebyshev's method.

You can find the exhaustive solution of this problem in Samarsky's book *The Theory of Difference Schema*, Nauka, Moscow, 1977 (in Russian)

6.2. Solving a Complete Problem of Eigenvalues by the Rotation Method

For the sake of simplicity, we shall first consider a real symmetric matrix A defined by equation (6.2). Note that the seeking of all the eigenvalues and eigenvectors of this matrix reduces to seeking an orthogonal matrix T such that the product

$$D = T'AT \quad (6.29)$$

is a diagonal matrix. Indeed, if such an orthogonal matrix T is found, the diagonal elements of the matrix D are the eigenvalues of the matrix A and the columns of the matrix T are the respective eigenvectors of the matrix A^* .

Let us introduce into consideration the spherical norm of the matrix A :

$$\|A\|_{\text{sph}} = \left[\sum_{i=1}^n \sum_{j=1}^n a_{ij}^2 \right]^{1/2}.$$

Then, evidently, the inequality

$$\sum_{i=1}^n a_{ii}^2 \leq \|A\|_{\text{sph}}^2 \quad (6.30)$$

is valid for the diagonal elements of the matrix A , and this inequality passes into an exact equality only in the case when the matrix A is diagonal.

It should be pointed out that *under the orthogonal transformation of the matrix A* (i.e. under the transformation of the form $\tilde{A} = UAR$, where U and R are orthogonal matrices), *the spherical norm of that matrix does not change**.

(* To prove this fact we denote the diagonal elements of the matrix D by $\lambda_1, \lambda_2, \dots, \lambda_n$ and set $e_k = \|e'_k\|$, where the elements e'_k of the column e_k satisfy the condition $e'_k = 0$ for $k \neq i$ and $e_k^k = 1$. Then, evidently. $De_k = \lambda_k e_k$, i.e. $T'ATE_k = \lambda_k e_k$ and since $T' = T^{-1}$, we have $ATE_k = \lambda_k T_k$. Consequently, Te_k are the eigenvectors of the matrix A .)

Hence it follows that the difference between transformation (6.29) and all the orthogonal transformations of the matrix A is that the former makes the sum of the squares of the diagonal elements of the transformed matrix maximal and the sum of the squares of all the off-diagonal elements of that matrix minimal.

The rotation method is an iterative method for which the matrix T indicated above can be found as the limit of the infinite product of elementary rotation matrices each of which has the form i

$$T_{jj}(\phi) = \left\| \begin{array}{cccc} 1 & & & 0 \\ & \ddots & & \\ & & 1 & \\ & & & \ddots \\ & & & & 1 & \cos \phi \dots & -\sin \phi \dots & 0 \\ & & & & & \vdots & \vdots & \\ & & & & & & 1 & \\ & & & & & & & \ddots \\ & & & & & & & & 1 & \sin \phi \dots & -\cos \phi \dots & 0 \\ & & & & & & & & & \vdots & \vdots & \\ & & & & & & & & & & 1 & \\ & & & & & & & & & & & 1 \end{array} \right\| \quad \begin{array}{l} \text{(the } i\text{th row)} \\ \\ \\ \text{(the } j\text{th row)} \end{array} \quad (6.31)$$

In the whole, the method of rotations consists in the construction of the sequence of matrices

$$A, A_1, A_2, \dots, A_v, A_{v+1}, \dots, \quad (6.32)$$

each successive matrix being obtained from the preceding matrix by means of the elementary step of the form $A_{v+1} = T'_{ij} A_v T_{ij}$.

If we simplify the notation by omitting the index v and consider one Step $\tilde{A} = T'_{ij} A T_{ij}$ performed with the aid of matrix (6.31), then we get the following expressions for the elements $\tilde{a}_{ij}$ of the transformed matrix A in terms of the elements a_{ij} of the matrix A

$$\begin{aligned}
\tilde{a} &= a_{ki} \text{ for } k \neq i, j, l \neq i, j, \\
\tilde{a}_{ij} &= a_{il} \cos \phi + a_{ji} \sin \phi \text{ for } l \neq i, j, \\
\tilde{a}_{ji} &= -a_{il} \sin \phi + a_{jl} \cos \phi \text{ for } l \neq i, j, \\
\tilde{a}_{li} &= a_{li} \cos \phi + a_{lj} \sin \phi \text{ for } l \neq i, j, \\
\tilde{a}_{lj} &= -a_{li} \sin \phi + a_{lj} \cos \phi \text{ for } l \neq i, j, \\
\tilde{a}_{ii} &= (a_{ii} \cos \phi + a_{ji} \sin \phi) \cos \phi + (a_{ij} \cos \phi + a_{jj} \sin \phi) \sin \phi. \\
\tilde{a}_{ji} &= (-a_{ii} \sin \phi + a_{ji} \cos \phi) \cos \phi + (-a_{ij} \sin \phi + a_{jj} \cos \phi) \sin \phi, \\
\tilde{a}_{jj} &= -(-a_{ii} \sin \phi + a_{ji} \cos \phi) \sin \phi + (-a_{ij} \sin \phi + a_{jj} \cos \phi) \cos \phi, \\
\tilde{a}_{ij} &= -(a_{ii} \cos \phi + a_{ji} \sin \phi) \sin \phi + (a_{ij} \cos \phi + a_{jj} \sin \phi) \cos \phi.
\end{aligned} \tag{6.33}$$

Relations (6.33) and the condition of symmetry of the matrix A yield the following equality that can be easily verified:

$$\begin{aligned}
\sum_{k=1}^n \sum_{\substack{l=1 \\ k \neq l}}^n a_{kl}^{-2} &= \sum_{k=1}^n \sum_{\substack{i=1 \\ k \neq l}}^n a_{kl}^2 - 2a_{ij}^2 \\
&+ \frac{1}{2} [(a_{jj} - a_{ii}) \sin 2\phi + 2a_{ij} \cos 2\phi]^2. \tag{6.34}
\end{aligned}$$

It follows from this equation that for the maximum reduction of the sum of the squares of all the off-diagonal elements, it is necessary to choose matrix (6.31) such that the following **two requirements** are met.

(1) it is required to choose the numbers i and j such that the square of the element a_{ij}^2 is the greatest among the squares of all the off-diagonal elements of the matrix A , i.e. to subordinate the choice of the numbers i and j to the condition

$$a_{ij}^2 = \max_{\substack{1 \leq k \leq n \\ 1 \leq l \leq n \\ k \neq l}} a_{kl}^2.$$

(2) the angle of rotation ϕ in the matrix (6.31) should be chosen such that the equality

$$(a_{jj} - a_{ii}) \sin 2\phi + 2a_{ij} \cos 2\phi = 0 \quad (6.35)$$

holds true.

Equation (6.35) defines uniquely the angle ϕ satisfying the conditions

$$\tan 2\phi = \frac{2a_{ij}}{a_{ii} - a_{jj}}, |\phi| \leq \frac{\pi}{4}. \quad (6.36)$$

This equation makes it possible to calculate $\cos \phi$ and $\sin \phi$ by the formulas

$$\cos \phi = \left\{ \frac{1}{2} [1 + (1 + p^2)^{-1/2}] \right\}^{1/2},$$

$$\sin \phi = \operatorname{sgn} p \left\{ \frac{1}{2} [1 - (1 + p^2)^{-1/2}] \right\}^{1/2},$$

where $p = \frac{2a_{ij}}{(a_{ii} - a_{jj})}$.

Note that if we have chosen matrix (6.31) such that the indicated requirements (1) and (2) are met, then equation (6.34) acquires the following form:

$$\sum_{k=1}^n \sum_{\substack{l=1 \\ k \neq l}}^n a_{kl}^{-2} = \sum_{k=1}^n \sum_{\substack{i=1 \\ k \neq l}}^n a_{kl}^2 - 2a_{ij}^2, \quad (6.37)$$

in which a_{ij} is an off-diagonal element of the matrix having the largest absolute value.

Now we can say more exactly that the rotation method consists constructing a sequence of matrices (6.32), each successive being obtained from the preceding matrix by means of the transformation $A_{v+1} = T'_{ij} A_v T_{ij}$ in which the matrix $T_{ij} = T_{ij}(\phi)$ is chosen such that the two requirements indicated above are met*.

Let us prove the convergence of the method of rotations. We denote by S_v^2 the sum of the squares of all the off-diagonal elements of the matrix A_v and by $a_{i_{vjv}}^{(v)}$ the off-diagonal element of this matrix which is the largest in the absolute value.

Then, by virtue of (6.37), there holds an equality

$$S_{v+1}^2 = S_v^2 - 2[a_{i_{vjv}}^{(v)}]^2. \quad (6.38)$$

Furthermore, since the total number of the off-diagonal elements of the matrix A is equal to $n(n-1)$, and $a_{i_{vjv}}^{(v)}$ is the greatest of these elements in absolute value, then the equality

$$[a_{i_{vjv}}^{(v)}]^2 \geq \frac{S_v^2}{n(n-1)} \quad (6.39)$$

holds true. But it follows from (6.38) and (6.39) that

$$S_{v+1}^2 \leq S_v^2 \left[1 - \frac{2}{n(n-1)} \right]. \quad (6.40)$$

Using successively inequality (6.40) written for the numbers $0, 1, \dots, v$ and designating as $S_0^2 = S_0^2(A)$ the sum of the squares of all the off-diagonal elements of the system matrix A , we find that

$$S_{v+1}^2 \leq S_0^2(A) \left[1 - \frac{2}{n(n-1)} \right]^{v+1}. \quad (6.41)$$

It immediately follows from inequality (6.41) that $\lim_{v \rightarrow \infty} S_{v+1}^2 = 0$, and this proves the convergence of the rotation method.

The diagonal elements of the matrix A_v are 'taken as approximate values of the eigen numbers of the matrix A and the columns of the matrix $T_{i_1 j_1} \cdot T_{i_2 j_2} \dots T_{i_v j_v}$ are taken as approximate eigenvectors of the matrix A .

A More exact results were obtained by Voevodin**. (**V.V. Voevodin, *Numerical Methods of Algebra. Theory and Algorithms*. Nauka, Moscow, 1966 (in Russian).)For the case when an arbitrary (not necessarily symmetric) matrix A does not have Jordan blocks and all its off-diagonal elements are values of order ε and are small as compared to the number $\rho = \lim_{\lambda_1 \neq \lambda_2} |\lambda_i - \lambda_j|$, obtained the following estimates:

(a) the estimate $\lambda_i = a_{ii} + \sum_{p=1}^n \frac{a_{ip} a_{pi}}{a_{ii} - a_{pp}} + O(\varepsilon^3)$ for the eigenvalues (we exclude,

from the indicated sum, the values p , belonging to the set R_j , of the numbers $j = 1, 2, \dots, n$ for which $\lambda_j = \lambda_i$);

(b) if T is a matrix whose columns are the eigenvectors of the matrix A and $T = I + H$, where I is an identity matrix, then the estimates

$$h_{ij} = \begin{cases} 0 & \text{if } \lambda_i = \lambda_j \\ \frac{a_{ij}}{a_{jj} - a_{ii}} + O(\varepsilon^3) & \text{if } \lambda_i \neq \lambda_j \end{cases}$$

are valid for the elements h_{ij} of the matrix H .

If A is a **complex Hermitian** matrix, then we must take the unitary matrix

$$T_{ij}(\phi, \psi) = \left\| \begin{array}{cccc} 1 & & & 0 \\ & \ddots & & \\ & & 1 & \\ & & \cos \phi & \dots \\ & \square & & -\sin \phi e^{-i\psi} \dots \\ & \square & 1 & \\ & \square & \vdots & \\ & \square & \ddots & 1 \\ \sin \phi e^{i\psi} & & \cos \phi & \dots \\ & & 1 & \ddots \\ & & & \ddots & 1 \\ 0 & & & & \end{array} \right\| \begin{array}{l} \text{(the } i\text{th row),} \\ \\ \text{(the } j\text{th row),} \end{array} \quad (6.42)$$

instead of matrix (6.31). Then, instead of equation (6.34), we arrive at the equation

$$\begin{aligned} \sum_{k=1}^n \sum_{\substack{l=1 \\ k \neq l}}^n |\tilde{a}_{kl}|^2 &= \sum_{k=1}^n \sum_{\substack{l=1 \\ k \neq l}}^n |a_{kl}|^2 - 2|a_{ij}|^2 \\ &+ 2|a_{ij}|^2 \cos^2 \phi - \sin^2 \phi e^{i\alpha} - \sin^2 \phi e^{i(2\psi - \alpha)} \\ &+ (a_{jj} - a_{ii}) \cos \phi \sin \phi e^{i\psi} \end{aligned}$$

in which α denotes the argument of the complex number a_{ij} .

To make the sum of the squares of the absolute values of the off-diagonal elements as small as possible, we must take the numbers i and j of matrix (6.42) such that the element a_{ij} is the greatest off-diagonal element of the matrix A in absolute value and subject the choice of the angles ϕ and ψ to the condition

$$(a_{jj} - a_{ii}) \cos \phi \cdot \sin \phi \cdot e^{i\psi} = 0.$$

The last condition leads to relations

$$\psi = \arg a_{ij}, \tan 2\phi = \frac{2|a_{ij}|}{a_{ii} - a_{jj}}, |\phi| \leq \frac{\pi}{4}.$$

The convergence of the rotation method can be proved in the same way as for the case of a real matrix.

Chapter – 7

BILINEAR AND QUADRATIC FORMS

In this chapter we study **bilinear forms** defined in a real linear space, i.e. the number functions of two vector arguments which are linear with respect to each argument. We study in detail the so-called quadratic forms which are **bilinear forms** defined for the coinciding values of their arguments. We also consider some applications of the theory of bilinear and quadratic forms.

7.1. Bilinear Forms

7.1.1. The concept of a bilinear form. The notion of a bilinear form in an arbitrary linear space was introduced in Chapter 5. However, to make the presentation of the material more convenient, we recall here some definitions and simple assertions.

Definition 1. *The number function $A(\mathbf{x}, \mathbf{y})$, whose arguments are various vectors $\mathbf{x}$ and $\mathbf{y}$ of the real linear space L , is called a **bilinear form** if the relations*

$$\left. \begin{aligned} A(\mathbf{x} + \mathbf{z}, \mathbf{y}) &= A(\mathbf{x}, \mathbf{y}) + A(\mathbf{z}, \mathbf{y}), \\ A(\mathbf{x}, \mathbf{y} + \mathbf{z}) &= A(\mathbf{x}, \mathbf{y}) + A(\mathbf{x}, \mathbf{z}), \\ A(\lambda \mathbf{x}, \mathbf{y}) &= \lambda A(\mathbf{x}, \mathbf{y}), \\ A(\mathbf{x}, \lambda \mathbf{y}) &= \lambda A(\mathbf{x}, \mathbf{y}). \end{aligned} \right\} \quad (7.1)$$

hold for any vectors $\mathbf{x}, \mathbf{y}$ and $\mathbf{z}$ from L and real number λ .

In other words, a bilinear form is a number function $A(\mathbf{x}, \mathbf{y})$ of two vector arguments $\mathbf{x}$ and $\mathbf{y}$ which is defined on various vectors $\mathbf{x}$ and $\mathbf{y}$ of the real linear space L and is linear with respect to each argument.*

(* It is often said in this case that the bilinear form $A(\mathbf{x}, \mathbf{y})$ is given on the linear space L .)

The simplest example of a bilinear form is the product of two linear forms $f(\mathbf{x})$ and $g(\mathbf{y})$ defined on the vectors $\mathbf{x}$ and $\mathbf{y}$ of the linear space L .

Definition 2. *The bilinear form $A(\mathbf{x}, \mathbf{y})$ is said to be **symmetric** (**skew-symmetric**) if the relations*

$$A(\mathbf{x}, \mathbf{y}) = A(\mathbf{y}, \mathbf{x}) \quad (A(\mathbf{x}, \mathbf{y}) = -A(\mathbf{y}, \mathbf{x})) \quad (7.2)$$

for any vectors $\mathbf{x}$ and $\mathbf{y}$ of the linear space L .

The following statement holds true: any bilinear form can be represented as the sum of a symmetric and a skew-symmetric bilinear form (see 5.9.1).

7.1.2. Representation of a bilinear form in a finite-dimensional linear space.

Suppose a bilinear form $B(\mathbf{x}, \mathbf{y})$ is given in an n -dimensional linear space L . Let us decide on the representation of the form $B(\mathbf{x}, \mathbf{y})$ in the case when a definite basis $e = (e_1, e_2, \dots, e_n)$ is specified in L .

The following assertion holds true.

Theorem 7.1. For the basis $e = (e_1, e_2, \dots, e_n)$ specified in an n -dimensional linear space, the bilinear form $B(\mathbf{x}, \mathbf{y})$ can be uniquely represented in the form

$$B(\mathbf{x}, \mathbf{y}) = \sum_{i, j=1}^n b_{ij} \xi_i \eta_j \quad (7.3)$$

where

$$b_{ij} = (e_i, e_j), \quad (7.4)$$

and ξ_i and η_j are the coordinates of the vectors $\mathbf{x}$ and $\mathbf{y}$, respectively, in the basis e .

Proof. Assume that $\mathbf{x} = \sum_{i=1}^n \xi_i e_i$ and $\mathbf{y} = \sum_{i=1}^n \eta_i e_i$ are the representations of the vectors $\mathbf{x}$ and $\mathbf{y}$ in terms of the basis e . The form $B(\mathbf{x}, \mathbf{y})$ being linear with respect to each of the arguments $\mathbf{x}$ and $\mathbf{y}$ (see (7.1)), we have

$$B(\mathbf{x}, \mathbf{y}) = B\left(\sum_{i=1}^n \xi_i e_i, \sum_{j=1}^n \eta_j e_j\right) = \sum_{i, j=1}^n B(e_i, e_j) \xi_i \eta_j.$$

Thus, representation (7.3) with expressions (7.4) for the coefficients b_{ij} is valid for the form $B(\mathbf{x}, \mathbf{y})$.

To prove the uniqueness of this representation, we assume that representation (7.3) with certain coefficients b_{ij} is valid for $B(x, y)$.

To prove the uniqueness of this representation, we assume that representation (7.3) with **certain** coefficients b_{ij} is valid for $B(x, y)$. Taking $x = e_i, y = e_j$ in (7.3), we immediately get expressions (7.4) for the coefficients b_{ij} . We have proved the theorem.

Definition. *The matrix*

$$B(e) = (b_{ij}) = \begin{pmatrix} b_{11} & b_{12} & \dots & b_{1n} \\ b_{21} & b_{22} & \dots & b_{2n} \\ \cdot & \cdot & \cdot & \cdot \\ b_{n1} & b_{n2} & \dots & b_{nn} \end{pmatrix}, \quad (7.5)$$

whose elements b_{ij} are defined with the aid of relations (7.4) called the **matrix of the bilinear form** $B(x, y)$ in the given basis e .

Remark 1. Let us consider the question of constructing all the bilinear forms in a given finite-dimensional real space L . The answer to this question is the following: any square matrix (b_{ij}) is a matrix of a certain bilinear form in the given basis $e = (e_1, e_2, \dots, e_n)$.

Let us verify the validity of this assertion.

Using the matrix (b_{ij}) , we shall define the number function $B(x, y)$ of two

arguments $x = \sum_{i=1}^n \xi_i e_i$ and $y = \sum_{j=1}^n \eta_j e_j$ of the form $B(x, y) = \sum_{i,j=1}^n b_{ij} \xi_i \eta_j$ in

the linear space L with the given basis $e = (e_1, e_2, \dots, e_n)$.

It is easy to see that this function satisfies all the conditions of the definition of a bilinear form. But then, according to Theorem 7.1, the elements b_{ij} of the given matrix are equal to $B(e_i, e_j)$, and the formula written above is the representation of that form as (7.3).

According to the remark made above, it is natural to call representation (7.3) of the bilinear form $B(x, y)$ the **general form of bilinear form in an n -dimensional linear space.** ;

Remark 2. If $B(x, y)$ is a symmetric (skew-symmetric) bilinear form, then matrix (7.5) of that form in the basis e is symmetric (skew-symmetric). The converse is also true: if matrix (7.5) of the bilinear form $B(x, y)$ is symmetric (skew-symmetric), then the linear form is symmetric (skew-symmetric) as well.

Let us verify the validity of this remark. §

Assume that $B(x, y)$ is a symmetric (skew-symmetric) bilinear form. Setting $x = e_i, y = e_j$ in (7.2), we find, in accordance with (7.4) that

$$b_{ij} = b_{ji}, (b_{ij} = -b_{ji}), \quad (7.6)$$

i.e. matrix (7.5) is symmetric (skew-symmetric).

Suppose now that matrix (7.5) of the bilinear form $B(x, y)$ is symmetric (skew-symmetric), i.e. its elements satisfy relations (7.6).

Then it follows from relation (7.3) and the relation $B(x, y) = \sum_{i,j=1}^n b_{ji} \xi_i \eta_j$ that B

$(x, y) = B(y, x), (B(x, y) = -B(y, x))$, i.e. the form $B(x, y)$ is symmetric (skew-symmetric).

7.1.3. Transformation of a matrix of a bilinear form in passing to anew basis.

The rank of a bilinear form. Let us consider two bases $e = (e_1, e_2, \dots, e_n)$ and $f = (f_1, f_2, \dots, f_n)$ in the linear space L . Assume that $A(e) = (a_{ij})$ and $A(f) = (b_{ij})$ are the matrices of the given bilinear form in the indicated bases.

Let us consider the connection between these matrices, i.e. the transformation of the matrix a_{ij} of the bilinear form when we pass from the basis e to a new basis f .

The following statement holds true.

Theorem 7.2. *The matrices $A(e)$ and $A(f)$ of the bilinear form $A(x, y)$ in the bases $e = (e_1, e_2, \dots, e_n)$ and $f = (f_1, f_2, \dots, f_n)$ are related as*

$$A(f) = C' A(e) C, \quad (7.7)$$

where $C = (c_{pq})$ is the transition matrix from the basis e to the basis f , and C' is the transposed matrix C .

Proof. The elements f_q of the new basis f can be expressed in terms of the elements e_p of the old basis e with the aid of the matrix $C = (c_{pq})$ by the formulas

$$f_q = \sum_{p=1}^n c_{pq} e_p. \quad (7.8)$$

Since $b_{lk} = A(f_l, f_k)$, we get, in accordance with (7.8),

$$\begin{aligned} x = e_p, y = b_{lk} = A(f_l, f_k) &= \left(\sum_{i=1}^n c_{ji} e_i, \sum_{j=1}^n c_{jk} e_j \right) \\ &= \sum_{i,j=1}^n A(e_i, e_j) c_{ij} c_{jk} = \sum_{i,j=1}^n a_{ij} c_{il} c_{jk}. \end{aligned} \quad (7.9)$$

Recall that the elements c'_{li} of the transposed matrix C' are connected with the elements c_{il} of the matrix C by the relation $c_{il} = c'_{li}$.

Substituting these relations into the right-hand side of (7.9), we get the following expression for b_{lk} ;

$$b_{lk} = \sum_{i,j=1}^n a_{ij} c'_{li} c_{jk} \sum_{i=1}^n c'_{li} \left(\sum_{j=1}^n a_{ij} c_{jk} \right). \quad (7.10)$$

The sum $\sum_{j=1}^n a_{ij} c_{jk}$ (by the definition of the product of matrices) is an element of the matrix $A(e)C$. Hence it follows that the expression on the right-hand side of (7.10) is an element of the matrix $C'A(e)C$. But the left-hand side of (7.10) includes an element of the matrix $A(f)$ and, therefore, $A(f) = C'A(e)C$. We have proved the theorem.

Corollary. *The rank of the matrix $A(f)$ is equal to that of the matrix $A(e)$.*

This follows immediately from (7.-7), from the fact that the matrix C , and, hence, the matrix C' are non-singular, and from the theorem stating that the rank of a matrix does not change when it is multiplied by a non-singular matrix.

This corollary makes it possible to introduce an important **numerical invariant** of a bilinear form, the so-called rank of a bilinear form.

Definition 1. *The rank of a bilinear form defined in a finite-dimensional linear space L is known as the rank of the matrix of that form in an arbitrary basis of the space L .*

Definition 2. *The bilinear form $A(x, y)$ given in the finite-dimensional linear space L is said to be non-singular (singular) if its rank is equal to (smaller than) the dimension of the space L .*

7.2. Quadratic Forms

Assume that $A(x, y)$ is a symmetric bilinear form defined on the linear space L .

Definition 1. *A quadratic form is a number function $A(x, y)$ of the same vector argument x which is obtained from the bilinear form $A(x, y)$ for $x = y$.*

The symmetric bilinear form $A(x, y)$ is said to be the polar of the quadratic form $A(x, x)$.

The polar bilinear form $A(x, y)$ and the quadratic form $A(x, x)$ are related as

$$A(x, y) = \frac{1}{2} [A(x + y, x + y) - A(x, x) - A(y, y)],$$

which follows from the obvious equality

$$A(x + y, x + y) = A(x, x) + A(x, y) + A(y, x) + A(y, y)$$

and the property of symmetry of the form $A(x, y)$.

Suppose a symmetric bilinear form $A(x, y)$ which is the polar of the quadratic form $A(x, x)$ is given in the finite-dimensional linear space L . Suppose, in addition, that a basis $e = (e_1, e_2, \dots, e_n)$ is specified in L .

In accordance with Theorem 7.1, the form $A(x, y)$ can be represented as (7.3)

$$A(\mathbf{x}, \mathbf{y}) = \sum_{i, j=1}^n a_{ij} \xi_i \eta_j, \quad (7.3)$$

where ξ_i and η_j are the coordinates of the vectors $\mathbf{x}$ and $\mathbf{y}$, respectively, in the basis e . By virtue of symmetry of $A(\mathbf{x}, \mathbf{y})$

$$a_{ij} = a_{ji} \quad (7.11)$$

(see Remark 2 in 7.1.2).

Setting $\mathbf{x} = \mathbf{y}$ (i.e. $\eta_j = \xi_j$) in (7.3), we get the following *representation for the quadratic form $A(\mathbf{x}, \mathbf{y})$ in the finite-dimensional space L with the given basis e* :

$$A(\mathbf{x}, \mathbf{x}) = \sum_{i, j=1}^n a_{ij} \xi_i \xi_j. \quad (1.12)$$

The matrix (a_{ij}) is called a **matrix of the quadratic form $A(\mathbf{x}, \mathbf{y})$ in the given basis e** .

According to (7.11), the matrix (a_{ij}) is symmetric. Evidently, every symmetric matrix (a_{ij}) is associated with the quadratic form $A(\mathbf{x}, \mathbf{x})$ with the aid of relation (7.12), relation (7.12) being a representation of $A(\mathbf{x}, \mathbf{x})$ in the space L with the given basis e (see also Remark 3 in 7.1.2).

Note that when we pass to a new basis, the matrix of a quadratic form is transformed by formula (7.7). Therefore, the rank of that matrix does not change when we pass to a new basis.

The rank of the matrix of the quadratic form $A(\mathbf{x}, \mathbf{x})$ is customarily called the **rank of a quadratic form**.

If the rank of the matrix of a quadratic form is equal to the dimension of the space L , the form is said to be **non-singular**, otherwise it is **singular**.

In what follows, we shall use the following terminology.

Definition 2. The quadratic form $A(\mathbf{x}, \mathbf{x})$ is called (1) **positive (negative) definite** if the inequality

$$A(\mathbf{x}, \mathbf{x}) > 0 \text{ (} A(\mathbf{x}, \mathbf{x}) < 0 \text{)}$$

is satisfied for any nonzero $\mathbf{x}$ (forms of this kind are also called **definite**);

(2) **alternating** if there are $\mathbf{x}$ and $\mathbf{y}$ such that

$$A(\mathbf{x}, \mathbf{x}) > 0, A(\mathbf{y}, \mathbf{y}) < 0;$$

(3) **quasi-definite** if

$$A(\mathbf{x}, \mathbf{x}) > 0 \text{ or } A(\mathbf{x}, \mathbf{x}) \leq 0$$

for all $\mathbf{x}$, but there is a nonzero $\mathbf{x}$ for which

$$A(\mathbf{x}, \mathbf{x}) = 0.$$

We shall later indicate the criteria which can be used to decide to which of the indicated types the form $A(\mathbf{x}, \mathbf{x})$ belongs.

Note the following important **assertion**.

If $A(\mathbf{x}, \mathbf{y})$ is a bilinear form which is the polar of the positive definite quadratic form $A(\mathbf{x}, \mathbf{x})$, then $A(\mathbf{x}, \mathbf{y})$ satisfies all the axioms of the scalar product of vectors in a Euclidean space.

Let us direct our attention to the four axioms of a scalar product (see 4.1.1).

If the number called the scalar product of the vectors $\mathbf{x}$ and $\mathbf{y}$ is symbolized as $A(\mathbf{x}, \mathbf{y})$, then the axioms can be written as follows?

$$1^\circ. A(\mathbf{x}, \mathbf{y}) = A(\mathbf{y}, \mathbf{x}).$$

$$2^\circ. A(\mathbf{x} + \mathbf{z}, \mathbf{y}) = A(\mathbf{x}, \mathbf{y}) + A(\mathbf{z}, \mathbf{y}).$$

$$3^\circ. A(\lambda\mathbf{x}, \mathbf{y}) = \lambda A(\mathbf{x}, \mathbf{y}).$$

$$4^\circ. A(\mathbf{x}, \mathbf{x}) \geq 0 \text{ and } A(\mathbf{x}, \mathbf{x}) > 0 \text{ for } \mathbf{x} \neq 0.$$

Since the bilinear form $A(\mathbf{x}, \mathbf{y})$, which is the polar of the quadratic form $A(\mathbf{x}, \mathbf{x})$, is symmetric, Axiom 1° is satisfied. Axioms 2° and 3°, in conjunction with the requirement of symmetry, are satisfied by virtue of the definition of a bilinear form (see 7.1.1). Axiom 4° is satisfied since the quadratic form $A(\mathbf{x}, \mathbf{x})$ is positive definite.

Remark. The axioms of a scalar product can evidently be regarded as a collection of requirements defining a bilinear form which is the polar of a positive definite quadratic form. A scalar product in linear spaces can, therefore, be defined with the aid of a bilinear form of this kind.

7.3. Reducing a Quadratic Form to the Sum of Squares

We describe here various methods of reducing a quadratic form to the sum of squares, i.e. show the ways of choosing the basis $f = (f_1, f_2, \dots, f_n)$ in the linear space L relative to which the quadratic form can be represented in the following **canonical form**:

$$A(x, x) = \lambda_1 \eta_1^2 + \lambda_2 \eta_2^2 + \dots + \lambda_n \eta_n^2, \quad (7.13)$$

$(\eta_1, \eta_2, \dots, \eta_n)$ being the coordinates of x in the basis f .

The coefficients $\lambda_1, \lambda_2, \dots, \lambda_n$ in expression (7.13) are called **canonical coefficients**.

It should be emphasized that we consider quadratic forms in an arbitrary real linear space. In 7.6 we shall study quadratic forms in a Euclidean space and shall prove the possibility of reducing every quadratic form to a canonical one even in an orthonormal basis. Proceeding from the results obtained in Chapter 5, we shall obtain, in 7.6, a new proof of the theorem on the reduction of a quadratic form to the canonical one in an arbitrary (not necessarily Euclidean) real linear space.

The present section is devoted not only to proving the possibility of reducing a quadratic form to the canonical form but also to a description of two methods of such a reduction which are very valuable in practical applications and are often encountered in practice.

Since every transformation of a basis is associated with a regular linear transformation of coordinates and a regular transformation of coordinates is associated with a transformation of a basis, the question concerning a reduction of a form to the canonical form can be decided by choosing a corresponding regular transformation of coordinates.

7.3.1. Lagrange's method. Let us prove the following theorem.

Theorem 7.3. *Any quadratic form $A(\mathbf{x}, \mathbf{x})$ defined in the n -dimensional linear space L can be reduced to the canonical form (7.13) by of a regular linear transformation of coordinates.*

Proof. We present the proof of the theorem by Lagrange's method. method consists in a successive completion of a quadratic trinomial with respect to every argument to form a perfect square.

We assume that $A(\mathbf{x}, \mathbf{x}) \neq 0$ (*If $A(\mathbf{x}, \mathbf{x}) \equiv 0$, then its matrix in any basis consists of zero elements and, therefore, such a form has a canonical form in any basis by definition.) and in the given basis $e = (e_1, e_2, \dots, e_n)$ has the form

$$A(\mathbf{x}, \mathbf{x}) = \sum_{i, j=1}^n a_{ij} \xi_i \xi_j. \quad (7.14)$$

We shall verify, first, that by means of a regular transformation of coordinates the form $A(\mathbf{x}, \mathbf{x})$ can be transformed so that the *coefficient in the square of the first coordinate of the vector $\mathbf{x}$ is nonzero*.

If in the given basis this coefficient is nonzero, then the required regular transformation is identical.

If $a_{11} = 0$, but the coefficient in the square of some other coordinate is different from zero, then we can renumber the base vectors and attain the necessary result. It is clear that the renumbering is a regular transformation.

Now if all the coefficients in the squares of the coordinates are zero, the necessary transformation can be obtained as follows. Assume, for instance, $a_{12} \neq 0$. Let us consider the following regular transformation of co-ordinates.

$$\xi'_1 = \xi_1 - \xi_2, \xi'_2 = \xi_1 + \xi_2, \xi'_i = \xi_i, i = 3, 4, \dots, n.$$

After this transformation, the coefficient in ξ'^2 is equal to $2a_{12}$ and is, therefore, nonzero.

Thus, we assume that $a_{11} \neq 0$ in relation (7.14).

Let us isolate, in expression (7.14), the group of the terms containing ξ_1 . We get,

$$A(x, x) = a_{11}\xi_1^2 + 2a_{12}\xi_1\xi_2 + \dots + 2a_{1n}\xi_1\xi_n + \sum_{i,j=1}^n a_{ij}\xi_i\xi_j. \quad (7.15)$$

We transform the isolated group of the terms as follows.

$$\begin{aligned} & a_{11}\xi_1^2 + 2a_{12}\xi_1\xi_2 + \dots + 2a_{1n}\xi_1\xi_n \\ &= a_{11} \left(\xi_1 + \frac{a_{12}}{a_{11}}\xi_2 + \dots + \frac{a_{1n}}{a_{11}}\xi_n \right)^2 \\ & \quad - \frac{a_{12}^2}{a_{11}}\xi_2^2 - \dots - \frac{a_{1n}^2}{a_{11}}\xi_n^2 - 2\frac{a_{12}a_{13}}{a_{11}}\xi_2\xi_3 \\ & \quad - \dots - 2\frac{a_{1n-1}a_{1n}}{a_{11}}\xi_{n-1}\xi_n. \end{aligned}$$

It is evident that expression (7.15) can be rewritten as

$$A(x, x) = a_{11} \left(\xi_1 + \frac{a_{12}}{a_{11}}\xi_2 + \dots + \frac{a_{1n}}{a_{11}}\xi_n \right)^2 + \sum_{i,j=1}^n a_{ij}^* \xi_i \xi_j, \quad (7.16)$$

where a_{ij}^* are coefficients in $\xi_i \xi_j$ Obtained as a result of the transformation.

Let us consider the following regular transformation of coordinates

$$\begin{aligned} \eta_1 &= \xi_1 + \frac{a_{12}}{a_{11}}\xi_2 + \dots + \frac{a_{1n}}{a_{11}}\xi_n, \\ \eta_2 &= \xi_2, \\ &\dots\dots\dots \\ \eta_n &= \xi_n. \end{aligned}$$

By means of this transformation and representation (7.16), for $A(x, x)$ we get

$$A(x, x) = a_{11}\eta_1^2 + \sum_{i,j=2}^n a_{ij}^* \eta_i \eta_j \cdot \xi_i \xi_j. \quad (7.17)$$

Thus, if the form $A(x, x) \neq 0$, then by means of regular transformation of coordinates this form can be reduced to (7.17).

Let us consider now the 'quadratic form $\sum_{i,j=2}^n a_{ij}^* \eta_i \eta_j$. If this form is identically zero, then the problem of reduction of $A(x, x)$ to the canonical form is solved.

Now if the form $\sum_{i,j=2}^n a_{ij}^* \eta_i \eta_j \neq 0$, can repeat the arguments, considering the transformation of the coordinates $\eta_2, \dots, \eta_n$, similar to those discussed above, without changing the coordinate η_1 . The transformation of coordinates $\eta_1, \eta_2, \dots, \eta_n$ of this type is, evidently, regular.

It is that in a finite number of steps we reduce the quadratic form $A(x, x)$ to the canonical form (7.13).

Note that we can obtain the necessary transformation of the original coordinates $\xi_1, \xi_2, \dots, \xi_n$ by multiplying the regular transformation obtained in the process of reasoning. We have thus proved the theorem.

Remark 1. The basis in which a quadratic form has a canonical form is called **canonical**. Note that a canonical basis is not uniquely defined.

Remark 2. If the form $A(x, x)$ is reduced to canonical equation (7.13), then, in general, not all the canonical coefficients λ_i are nonzero. Leaving only nonzero coefficients λ_i in (7.13) and renumbering them, we get the following expression for $A(x, x)$:

$$A(x, x) = \lambda_1 \eta_1^2 + \lambda_2 \eta_2^2 + \dots + \lambda_k \eta_k^2. \quad (7.18)$$

It is clear that $k \leq n$. Since, by definition, the rank of a quadratic form is equal to the rank of its matrix in any basis, it follows from (7.18) and the condition $\lambda_i \neq 0$

for $i = 1, 2, \dots, k$ that the rank of the form is k . Thus, *the number of nonzero canonical coefficients is equal to the rank of the quadratic form.*

7.3.2. Jacobi's method. Under some additional assumptions concerning the quadratic form $A(x, x)$, we can indicate explicit formulas of transition from the given basis $e = (e_1, e_2, \dots, e_n)$ to the canonical basis $f = (f_1, f_2, \dots, f_n)$ and formulas for the canonical coefficients λ_i :

We first introduce the concept of the **triangular transformation** of the base vectors.

The transformation of the base vectors $e_1, e_2, \dots, e_n$ is said to be triangular if it has the form

$$\left. \begin{aligned} f_1 &= e_1, \\ f_2 &= \alpha_{21}e_1 + e_2, \\ f_3 &= \alpha_{31}e_1 + \alpha_{32}e_2, \\ &\dots\dots\dots \\ f_n &= \alpha_{n1}e_1 + \alpha_{n2}e_2 + \dots + e_n. \end{aligned} \right\} \quad (7.19)$$

Remark. Since the determinant of the matrix of the triangular transformation (7.19) is nonzero (equal to 1), the vectors $f_1, f_2, \dots, f_n$ form a basis.

We introduce into consideration the angular minors of the matrix $A(e) = (a_{ij})$ of the coefficients of the form $A(x, x)$ in the basis e , designating them as $\Delta_1, \Delta_2, \dots, \Delta_{n-1}$:

$$\Delta_1 = a_{11}, \Delta_2 = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix}, \dots, \Delta_{n-1} = \begin{vmatrix} a_{11} & \dots & a_{1, n-1} \\ \dots\dots\dots & & \\ a_{n-1, 1} & \dots & a_{n-1, n-1} \end{vmatrix}. \quad (7.20)$$

The following statement holds true.

Theorem 7.4. *Let the minors $\Delta_1, \Delta_2, \dots, \Delta_{n-1}$ of the matrix (a_{ij}) of the quadratic form $A(x, x)$ be nonzero. Then there is a unique triangular transformation of the base*

vectors $e_1, e_2, \dots, e_n$ by means of which the form $A(x, x)$ can be reduced to the canonical form.

Proof. Recall that the coefficients b_{ij} of the form $A(x, x)$ in the basis $f_1, f_2, \dots, f_n$ can be calculated by the formulas $b_{ij} = A(f_i, f_j)$.

If in the basis $f_1, f_2, \dots, f_n$ $A(x, x)$ has a canonical form, then $b_{ij} = 0$ for $i \neq j$. Therefore, to prove the theorem it is sufficient to construct, by means of the triangular transformation (7.19), a basis $f_1, f_2, \dots, f_n$ such that the following relations hold true :

$A(f_i, f_j) = 0$ for $i \neq j$, or, which is the same, for $i < j$ (7.21) (we must verify, of course, that the required transformation is unique).

If we direct our attention to formulas (7.19) for f_i , then, using the linear property of the quadratic form $A(x, x)$ with respect to each argument, we can easily see that relations (7.21) are satisfied if the relations .

$$A(e_1, f_j) = 0, A(e_2, f_j) = 0, \dots, A(e_{j-1}, f_j) = 0, \quad (7.22)$$

$$j = 2, 3, \dots, n$$

hold true.

Let us write out formulas (7.22). For that purpose, we substitute the expression

$$f_j = \alpha_{j1}e_1 + \alpha_{j2}e_2 + \dots + \alpha_{j, j-1}e_{j-1} + e_j \quad (7.23)$$

for f_j from relations (7.19) into the left-hand sides of these formulas. Using the linear property of $A(x, x)$ with respect to each argument and the designation $A(e_i, e_j) = a_{ij}$, we obtain the following linear system of equations for the unknown coefficients α_{jk} :

$$\left. \begin{aligned} \alpha_{j1}a_{11} + \alpha_{j2}a_{12} + \dots + \alpha_{j, j-1}a_{1, j-1} + a_{1j} &= 0, \\ \alpha_{j1}a_{j-1, 1} + \alpha_{j2}a_{j-1, 2} + \dots + \alpha_{j, j-1}a_{j-1, j-1} + a_{j-1, j} &= 0. \end{aligned} \right\} \quad (7.24)$$

The determinant of this system is equal to Δ_{j-1} . By the hypothesis, $\Delta_{j-1} \neq 0$. Consequently, system (7.24) has a unique solution. We can thus construct a unique triangular transformation of the base vectors by means of which $A(x, x)$ can be reduced to the canonical form. We have proved the theorem.

We present formulas for calculating the coefficients α_{ji} of the required triangular transformation and the formulas for the canonical coefficients Δ_j .

We use the symbol $\Delta_{j-1, i}$ for the/minor of the matrix (a_{ij}) lying in the rows of the matrix labelled $1, 2, \dots, j-1$ and the columns labelled $1, 2, \dots, i-1, i+1, \dots, j$. Then, using system (7.24) and Cramer's rule, we get the following expression for α_{ji} :

$$\alpha_{ji} = (-1)^{j+i} \frac{\Delta_{j-1, i}}{\Delta_{j-1}}. \quad (7.25)$$

Let us now calculate the canonical coefficients λ_j .

Since $\lambda_j = b_{jj} = A(f_j, f_j)$, expression (7.23) for f_j and formulas (7.22) yield

$$\begin{aligned} \lambda_j = A(f_j, f_j) &= A(\alpha_{j1}e_2 + \alpha_{j2}e_2 + \dots + \alpha_{j, j-1}e_{j-1} \\ &+ e_j, f_j) = A(e_j, \alpha_j f_j) = A(e_j, \alpha_{j1}e_1 + \alpha_{j2}e_2 \\ &+ \dots + \alpha_{j, j-1}e_{j-1} + e_j) = \alpha_{j1}a_{1j} + \alpha_{j2}a_{2j} \\ &+ \dots + \alpha_{j, j-1}a_{j-1, j} + a_{jj}. \end{aligned}$$

Substituting expression (7.25) for α_{ji} into the right-hand side of the last relation, we find that

$$\begin{aligned} \lambda_j &= ((-1)^{j+1} a_{1j} \Delta_{j-1, 1} + (-1)^{j+2} a_{2j} \Delta_{j-1, 2} + \dots \\ &\dots + (-1)^{2j-1} a_{j-1, j} \Delta_{j-1, j-1} + a_{jj} \Delta_{j-1, j-1}) \Delta_{j-1}^{-1}. \end{aligned}$$

The numerator of the last relation is the sum of the products of the elements of the row labelled j in the determinant Δ_j by the algebraic adjuncts of those elements in the indicated determinant. Consequently, the numerator is equal to Δ_j . Therefore,

$$\lambda_j = \frac{\Delta_j}{\Delta_{j-1}}, j = 2, 3, \dots, n. \quad (7.26)$$

Since $\lambda_1 = A(f_1, f_1) = A(e_1, e_1) = a_{11} = \Delta_1$, we get from this and from (7.26) the following formulas for the canonical coefficients;

$$\lambda_1 = \Delta_1, \lambda_2 = \frac{\Delta_2}{\Delta_1}, \dots, \lambda_n = \frac{\Delta_n}{\Delta_{n-1}}. \quad (7.27)$$

7.4. The Law of Inertia of Quadratic Forms. Classification of Quadratic Forms .

7.4.1. The law of inertia of quadratic forms. We have pointed out (see Remark 2 in 7.3.1) that the rank of a quadratic form is equal to the number of nonzero canonical coefficients. Thus, the number of nonzero canonical coefficients does not depend on the choice of a regular transformation by means of which the form $A(x, x)$ can be reduced to the canonical form. Indeed, the number of positive and negative canonical coefficients does not change with any method of reducing $A(x, x)$ to the canonical form. This property is called the **law of inertia** of quadratic forms.

Some remarks are due before passing to the substantiation of the law of inertia.

Suppose in the basis $e = (e_1, e_2, \dots, e_n)$ the form $A(x, x)$ is defined by the matrix $A(e) = (a_{ij})$:

$$A(x, x) = \sum_{i, j=2}^n a_{ij} \xi_i \xi_j, \quad (7.28)$$

where $\xi_1, \xi_2, \dots, \xi_n$ are the coordinates of the vector x in the basis e . We assume that with the aid of regular transformation of coordinates, this form has been reduced to the canonical form

$$A(\mathbf{x}, \mathbf{x}) = \lambda_1 \mu_1^2 + \lambda_2 \mu_2^2 + \dots + \lambda_k \mu_k^2, \quad (7.29)$$

$\lambda_1, \lambda_2, \dots, \lambda_k$, being nonzero canonical coefficients labelled so that the first q coefficients are positive and the following coefficients are negative:

$$\lambda_1 > 0, \lambda_2 > 0, \dots, \lambda_q > 0, \lambda_{q+1} < 0, \dots, \lambda_k < 0.$$

Let us consider the following regular transformation of the coordinates μ_i^* :

$$\left. \begin{aligned} \eta_1 &= \frac{1}{\sqrt{\lambda_1}} \mu_1, \eta_2 = \frac{1}{\sqrt{\lambda_2}} \mu_2, \dots, \eta_q = \frac{1}{\sqrt{\lambda_q}} \mu_q, \\ \eta_{q+1} &= \frac{1}{\sqrt{-\lambda_{q+1}}} \mu_{q+1}, \dots, \eta_k = \frac{1}{\sqrt{-\lambda_k}} \mu_k, \\ \eta_{k+1} &= \mu_{k+1}, \dots, \eta_n = \mu_n. \end{aligned} \right\} \quad (7.30)$$

As a result of this transformation $A(\mathbf{x}, \mathbf{x})$ assumes the form

$$A(\mathbf{x}, \mathbf{x}) = \eta_1^2 + \eta_2^2 + \dots + \eta_q^2 - \eta_{q+1}^2 - \dots - \eta_k^2, \quad (7.31)$$

which is called the *normal form* of the quadratic form.

Thus, by means of a certain regular transformation of the coordinates $\xi_1, \xi_2, \dots, \xi_n$ of the vector $\mathbf{x}$ in the basis $\mathbf{e} = (\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n)$

$$\begin{aligned} \eta_i &= \alpha_{i1} \xi_1 + \alpha_{i2} \xi_2 + \dots + \alpha_{in} \xi_n, \\ i &= 1, 2, \dots, n, \det(\alpha_{ij}) \neq 0. \end{aligned} \quad (7.32)$$

(this transformation is the product of the transformations of ξ into μ and μ into η formulas (7.30)), the quadratic form can be reduced normal form (7.31).

Let us prove the following assertion.

Theorem 7.5 (the law of inertia of quadratic forms). *The number of terms with positive (negative) coefficients in the normal form of a quadratic form does not depend on the way of reducing the quadratic form to that form.*

Proof. Assume that by means of regular transformation of coordinates (7.32) the form $A(x, x)$ has been reduced to the normal form (7.31) and by means of another regular transformation of coordinates it has been reduced to the normal form

$$A(x, x) = \xi_1^2 + \xi_2^2 + \dots + \xi_p^2 - \xi_{p+1}^2 - \dots - \xi_k^2. \quad (7.33)$$

To prove the theorem, it is, evidently, sufficient to verify the validity of the equality $p = q$.

Suppose $p > q$. Let us verify that in this case there is a non-zero vector x such that with respect to the bases in which this form $A(x, x)$ has the forms (7.31) and (7.33), the coordinates $\eta_1, \eta_2, \dots, \eta_q$ and $\xi_{p+1}, \dots, \xi_q$ of that vector are equal to zero :

$$\eta_1 = 0, \eta_2 = 0, \dots, \eta_l = 0, \xi_{p+1} = 0, \dots, \xi_n = 0. \quad (7.34)$$

Since the coordinates η_i have been obtained by means of regular transformation (7.32) of the coordinates $\xi_1, \dots, \xi_n$, and the coordinates ξ_j have been obtained by means of a similar regular transformation of the same coordinates ξ_1, ξ_n , relations (7.34) can be regarded as a system of homogeneous linear equations with respect to the coordinates $\xi_1, \dots, \xi_n$ of the required vector x in the basis $e = (e_1, e_2, \dots, e_n)$ (for example, in accordance with (7.32), the relation $\eta_1 = 0$ can be written out as $\alpha_{11}\xi_1 + \alpha_{12}\xi_2 + \dots + \alpha_{1n}\xi_n = 0$). Since $p > q$, the number of homogeneous equations (7.34) is less than n and, therefore, system (7.34) has a nonzero solution with respect to the coordinates $\xi_1, \dots, \xi_n$ of the required vector x . Consequently, if $p > q$, then there is a nonzero vector x such that relations (7.34) hold true.

Let us calculate the value of the form $A(x, x)$ for the vector x . Using relations (7.31) and (7.33), we get

$$A(x, x) = -\eta_{q+1}^2 - \dots - \eta_k^2 = \xi_1^2 + \xi_2^2 + \dots + \xi_p^2.$$

The last equality holds only in the case $\eta_{q+1} = \dots = \eta_k = 0$ and $\xi_1 = \xi_2 = \dots = \xi_p = 0$. Thus, in a certain basis all the coordinates $\xi_1, \xi_2, \dots, \xi_n$ of the nonzero vector $\mathbf{x}$ are equal to zero (see the last equations and relations (7.34)), i.e. the vector $\mathbf{x}$ is equal to zero. Consequently, the assumption $p > q$ leads to a contradiction. By similar considerations $p < q$ leads to a contradiction too.

Thus, $p = q$. We have proved the theorem.

7.4.2. Classification of quadratic forms. In 7.2.1 (see Definition 2) we introduced the concepts of a positive definite, negative definite, alternating and quasi-definite quadratic forms.

In this subsection, with the aid of the concepts of inertia indices, of positive and negative indices of inertia of a quadratic form, we shall indicate the way of determining the type of a quadratic form. The *index of inertia* of a quadratic form is the number of non-zero canonical coefficients of that form (i.e. its rank), the *positive index of inertia* is the number of positive canonical coefficients and the *negative index of inertia* is the number of negative canonical coefficients. It is clear that the sum of a positive and a negative index of inertia is equal to the index of inertia.

Thus, assume that the index of inertia, the positive and the negative index of inertia of the quadratic form $A(\mathbf{x}, \mathbf{x})$ are equal to k, p and q respectively ($k = p + q$). We proved in 7.4.1 that in any canonical basis $\mathbf{f} = (f_1, f_2, \dots, f_n)$ this form could be reduced to the following normal form:

$$A(\mathbf{x}, \mathbf{x}) = \eta_1^2 + \eta_2^2 + \dots + \eta_p^2 - \eta_{p+1}^2 - \dots - \eta_k^2, \quad (7.35)$$

where $\eta_1, \eta_2, \dots, \eta_n$ are the coordinates of the vector $\mathbf{x}$ in the basis $\mathbf{f}$.

1°. *The necessary and sufficient condition for a quadratic form to be definite.*

The following statement holds true.

For the quadratic form $A(\mathbf{x}, \mathbf{x})$ given in the n -dimensional linear space L to be definite, it is necessary and sufficient that either the positive index of inertia p or the negative index of inertia q , should be equal to the dimension n of the space L .

In this case, if $p = n$, then the form is positive definite, if $q = n$, then the form is negative definite.

Proof. Since the cases of a positive definite form and a negative definite form are treated analogously, we shall give the proof of the statement for positive definite forms.

(1) Necessity. Let the form $A(x, x)$ be positive definite. Then expression (7.35) assumes the form

$$A(x, x) = \eta_1^2 + \eta_2^2 + \dots + \eta_p^2.$$

If in this case $p < n$, then it follows from the last expression that for the nonzero vector x with the coordinates

$$\eta_1 = 0, \eta_2 = 0, \dots, \eta_p = 0, \eta_{p+1} \neq 0, \dots, \eta_n \neq 0$$

the form $A(x, x)$ vanishes, and this contradicts the definition of a positive definite quadratic form. Consequently, $p = n$.

(2) Sufficiency. Assume $p = n$. Then relation (7.35) has the form $A(x, x) = \eta_1^2 + \eta_2^2 + \dots + \eta_n^2$. It is clear that $A(x, x) \geq 0$, and if $A = 0$, then $\eta_1 = \eta_2 = \dots = \eta_n = 0$, i.e. the vector x is zero. Consequently, $A(x, x)$ is a positive definite form.

Remark. To use the indicated criterion to decide whether a quadratic form is definite, we must reduce the quadratic form to the canonical form.

In 7.4.3 we shall prove *Sylvester's criterion* for the definiteness of a quadratic form which can be used to decide whether a form given in any basis without being reduced to a canonical form is definite.

2°. *The necessary and sufficient condition for a quadratic form to be alternating.*

Let us prove the following statement.

For a quadratic form to be alternating, it is necessary and sufficient that both the positive and the negative indices of inertia of that form be nonzero.

Proof. **(1) Necessity.** Since an alternating form can assume positive as well as negative values, its representation (7.35) in the normal form must contain both positive and negative terms (otherwise, the form would assume either nonnegative or non-positive values). Consequently, both positive and negative indices of inertia are different from zero.

(2) Sufficiency. Assume $p \neq 0$ and $q \neq 0$. Then we have $A(\mathbf{x}_1, \mathbf{x}_1) > 0$ for the vector $\mathbf{x}_1$ with the coordinates $\eta_1 \neq 0, \dots, \eta_p \neq 0, \eta_{p+1} = 0, \dots, \eta_n = 0$ and $A(\mathbf{x}_2, \mathbf{x}_2) < 0$ for the vector $\mathbf{x}_2$ with the coordinates $\eta_1 = 0, \dots, \eta_p = 0, \eta_{p+1} \neq 0, \dots, \eta_n \neq 0$. Consequently, $A(\mathbf{x}, \mathbf{x})$ is an alternating form.

3°. *The necessary and sufficient condition for a quadratic form to be quasi-definite.*

The following statement holds true.

For the form $A(\mathbf{x}, \mathbf{x})$ to be quasi-definite, it is necessary and sufficient that the following relations hold true: either $p < n, q = 0$ or $p = 0, q < n$.

Proof. We shall consider the case of a positive quasi-definite form. The case of a negative quasi-definite form can be treated by analogy.

(1) Necessity. Let the form $A(\mathbf{x}, \mathbf{x})$ be positive quasi-definite. Then, evidently, $q = 0$ and $p < n$ (in the case $p = n$ the form would be positive definite).

(2) Sufficiency. If $p < n, q = 0$, then $A(\mathbf{x}, \mathbf{x}) \geq 0$ and we have $A(\mathbf{x}, \mathbf{x}) = 0$ for the nonzero vector $\mathbf{x}$ with the coordinates $\eta_1 = 0, \dots, \eta_p = 0, \eta_{p+1} \neq 0, \dots, \eta_n \neq 0$. i.e. $A(\mathbf{x}, \mathbf{x})$ is a positive quasi-definite form.

7.4.3. Sylvester's* criterion for the definiteness of a quadratic form. Assume that the form $A(\mathbf{x}, \mathbf{x})$ in the basis $e = (e_1, e_2, \dots, e_n)$ is defined by the matrix

$$A(e) = (a_{ij}): A(\mathbf{x}, \mathbf{x}) = \sum_{i,j=1}^n a_{ij} \xi_i \xi_j, \quad \text{and} \quad \text{that}$$

determinant of the matrix (a_{ij}) . The following statement holds true.

Theorem 7.6 (Sylvester's criterion). *For the quadratic form $A(x, x)$ to be positive definite, it is necessary and sufficient that the inequalities $\Delta_1 > 0, \Delta_2 > 0, \dots, \Delta_n > 0$ be satisfied.*

Proof. (I) Necessity. We shall first prove that the condition of definiteness of the quadratic form $A(x, x)$ yields $\Delta_i \neq 0, i=1, 2, \dots, n$.

Thus, assume $\Delta_k = 0$. Let us consider the following quadratic homogeneous system of linear equations:

Since Δ_k is a determinant of this system and $\Delta_k \neq 0$, system (7.36) has a nonzero solution $\xi_1, \xi_2, \dots, \xi_k$ (not all ξ_i are zero). Let us multiply the first equation (7.36) by ξ_1 , the second equation by $\xi_2, \dots$, the last equation by ξ_k and add the resulting relations together. As a result we obtain an equality $\sum_{i,j=1}^n a_{ij} \xi_i \xi_j = 0$, whose left-

hand side is a value of the quadratic form $A(x, x)$ for the nonzero vector x with the coordinates $(\xi_1, \xi_2, \dots, \xi_k, 0, \dots, 0)$. This value is equal to zero which contradicts the definiteness of the system.

We have thus verified that $\Delta_i \neq 0$, $i = 1, 2, \dots, n$. We can, therefore, apply Jacobi's method for reducing the form $A(x, x)$ to the sum of squares (see Theorem 7.4) and use formulas (7.27) for the canonical coefficients λ_i . If $A(x, x)$ is a positive definite form, then all the canonical coefficients are positive. But then it follows from relations (7.27) that $\Delta_1 > 0, \Delta_2 > 0, \dots, \Delta_n > 0$. Now if $A(x, x)$ is a negative definite form, then all the canonical coefficients are negative. But then it follows from formulas (7.27) that the signs of the angular minors alternate, and $\Delta_1 < 0$.

(2) Sufficiency. Suppose the conditions imposed on the angular minors Δ_i in the formulation of the theorem are fulfilled. Since $\Delta_i \neq 0$, $i = 1, 2, \dots, n$, the form A can be reduced to the sum of squares by Jacobi's method (see Theorem 7.4), and the canonical coefficients λ_i can be found by formulas (7.27). $\Delta_1 > 0, \Delta_2 > 0, \dots, \Delta_n > 0$, then it follows from relations (7.27) that all $\lambda_i > 0$, i.e. the form $A(x, x)$ is positive definite. Now if the signs of Δ_i alternate and $\Delta_1 < 0$, then it follows from relations (7.27) that the form $A(x, x)$ is negative definite. And this is what we wished to prove.

7.5. Multilinear Forms

Definition. The multilinear form $A(x_1, x_2, \dots, x_p)$ of p vector arguments is a number function which is defined on various vectors $x_1, x_2, \dots, x_p$ of the linear space L and is linear with respect to each argument, the values of the other arguments being fixed.

The product of the linear forms $A(x_1)A(x_2), \dots, A(x_p)$ can serve as a simple example of a multilinear form.

The multilinear form $A(x_1, x_2, \dots, x_p)$ is said to be *symmetric* (*skew-symmetric*) if for every two of its arguments x_k and x_l for any values of these arguments the following relation holds true :

$$A(x_1, \dots, x_k, \dots, x_l, \dots, x_p) = A(x_1, \dots, x_l, \dots, x_k, \dots, x_p)$$

$$(A(x_1, \dots, x_k, \dots, x_l, \dots, x_p) = -A(x_1, \dots, x_l, \dots, x_k, \dots, x_p)).$$

Suppose the multilinear form $A(x_1, x_2, \dots, x_p)$ is defined in the finite-dimensional linear space L and let $e_1, e_2, \dots, e_n$ be a basis in L . Let us resolve each vector x_i into the components $e_1, e_2, \dots, e_n$:

Substituting the expressions for x_i formulas (7.37) into the multilinear form $A(x_1, x_2, \dots, x_p)$ and using the property of linearity of that form with respect to each argument, we get

$$A(x_1, x_2, \dots, x_p) = A\left(\sum_{j_1=1}^n \xi_{1j_1} e_{j_1}, \sum_{j_2=1}^n \xi_{2j_2} e_{j_2}, \dots, \sum_{j_p=1}^n \xi_{pj_p} e_{j_p}\right)$$

$$= \sum_{j_1, j_2, \dots, j_p=1}^n \xi_{1j_1} \xi_{2j_2} \dots \xi_{pj_p} A(e_{j_1}, e_{j_2}, \dots, e_{j_p}). \quad (7.38)$$

Thus the values of the multilinear form $A(x_1, x_2, \dots, x_p)$ in a finite-dimensional space with the isolated basis $e_1, e_2, \dots, e_n$ are define by various values $A(e_{j_1}, e_{j_2}, \dots, e_{j_p})$ of that form on the vectors $e_{j_1}, e_{j_2}, \dots, e_{j_p}$.

Let us prove the following assertion.

Theorem 7.7. *Any multilinear skew-symmetric form $A(x_1, x_2, \dots, x_n)$ defined in the n -dimensional linear space L with the isolated basis $e_1, e_2, \dots, e_n$ can be represented as*

$$A(x_1, x_2, \dots, x_n) = \alpha \begin{vmatrix} \xi_{11} & \xi_{12} & \dots & \xi_{1n} \\ \xi_{21} & \xi_{22} & \dots & \xi_{2n} \\ \dots & \dots & \dots & \dots \\ \xi_{n1} & \xi_{n2} & \dots & \xi_{nn} \end{vmatrix} \quad (7.39)$$

where $a = A(e_1, e_2, \dots, e_n)$, $a(\xi_{i1}, \xi_{i2}, \dots, \xi_{in})$ are the coordinates of the vectors x_i in the basis $e_1, e_2, \dots, e_n$.

Proof. Since the form $A(x_1, x_2, \dots, x_n)$ is skew-symmetric to perform an arbitrary permutation $(j_1, j_2, \dots, j_n)$ of the indices $(1, 2, \dots, n)$ we have

$$\begin{aligned} A(e_{j_1}, e_{j_2}, \dots, e_{j_n}) &= (-1)^{N(j_1, j_2, \dots, j_n)} A(e_1, e_2, \dots, e_n) \\ &= (-1)^{N(j_1, j_2, \dots, j_n)} a, \end{aligned} \quad (7.40)$$

where $N(j_1, j_2, \dots, j_n)$ is the number of disorders in the permutation $(j_1, j_2, \dots, j_n)$.

Because of the skew-symmetry of the form, the value $A(e_{j_1}, \dots, e_{j_k}, \dots, e_{j_l}, \dots, e_{j_n})$ is equal to zero for two identical indices j_k and j_l ($j_k = j_l$). It follows from this and from (7.40) that for the case being considered relation (7.38) assumes the form

$$\begin{aligned} &A(x_1, x_2, \dots, x_n) \\ &= \sum_{j_1, j_2, \dots, j_n=1}^n (-1)^{N(j_1, j_2, \dots, j_n)} \xi_{1j_1} \xi_{2j_2} \dots \xi_{nj_n}. \end{aligned} \quad (7.38)$$

Comparing formula (7.41) with formula (1.28) in Chapter I for a determinant of order n , we ascertain the validity of relation (7.39). We have proved the theorem.

7.6. Bilinear and Quadratic Forms in a Euclidean Space

In previous sections we studied bilinear and quadratic forms in an arbitrary (not necessarily Euclidean) real linear space L . Here we shall obtain data on bilinear and quadratic forms defined in a real Euclidean space. We shall widely use the results obtained in 5.9 devoted to linear operators.

It will be shown in 7.6.3 how the theory of Euclidean spaces can be applied to obtain comprehensive results in arbitrary linear spaces. In particular, we shall obtain an independent proof of the theorem stating that every quadratic form in a linear space can be reduced to a canonical form.

7.6.1. Preliminary remarks. We shall recall here some notions of the theory of linear operators.

Let V be a n -dimensional real Euclidean space and A , a linear operator from V into V . The operator A^* is said to be an adjoint of A if the equality

$$(Ax, y) = (x, A^*y) \quad (7.42)$$

is valid for all $x \in V$ and $y \in V$.

The operator A is said to be self-adjoint if $A = A^*$, i.e. if

$$(Ax, y) = (x, Ay) \quad (7.43)$$

for all $x \in V$ and $y \in V$.

Let us consider a bilinear form $B(x, y)$ defined in the Euclidean space V . We have found (see Chapter 5) that every such form $B(x, y)$ is in a one-to-one correspondence with a linear operator such that the equality

$$B(x, y) = (Ax, y) \quad (7.44)$$

holds true.

In addition, it has been proved (Theorem 5.33) that the bilinear form $B(x, y)$ is symmetric if and only if the operator A appearing in (7.44) is self-adjoint.

It should also be pointed out that it has been proved (Theorem 5.35) that any self-adjoint operator A has an orthonormal basis of eigenvectors. This means that there

exists an orthonormal system $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ and real numbers $\lambda_1, \lambda_2, \dots, \lambda_n$ such that

$$\mathbf{A}\mathbf{e}_k = \lambda_k \mathbf{e}_k. \quad (7.45)$$

Note that in the basis $\{\mathbf{e}_k\}$ the matrix of the operator $\mathbf{A}$ has a diagonal form.

7.6.2. Reducing a quadratic form to the sum of squares in an orthogonal basis.

Suppose $B(\mathbf{x}, \mathbf{y})$ is a symmetric bilinear form defined in the real Euclidean space V , and $B(\mathbf{x}, \mathbf{x})$ is a quadratic form it defines.

Let us prove the following theorem on reducing the quadratic form $B(\mathbf{x}, \mathbf{x})$ to the sum of squares.

Theorem 7.8. *Let $B(\mathbf{x}, \mathbf{y})$ be a symmetric bilinear form defined in the Euclidean space V . Then there exists, in the space V , an orthonormal basis $\{\mathbf{e}_k\}$ and real numbers such that for any $\mathbf{x} \in V$ the quadratic form $B(\mathbf{x}, \mathbf{x})$ can be represented as the following sum of the squares of the coordinates of the vector $\mathbf{x}$ in the basis $\{\mathbf{e}_k\}$:*

$$B(\mathbf{x}, \mathbf{x}) = \sum_{k=1}^n \lambda_k \xi_k^2. \quad (7.46)$$

Proof. Since $B(\mathbf{x}, \mathbf{y})$ is a symmetric bilinear form, there is a self-adjoint operator $\mathbf{A}$ such that

$$B(\mathbf{x}, \mathbf{y}) = (\mathbf{A}\mathbf{x}, \mathbf{y}). \quad (7.47)$$

By Theorem 5.35, we can indicate, for the operator $\mathbf{A}$, an orthonormal basis $\{\mathbf{e}_k\}$ of the eigenvectors of the operator; let λ_k be the eigen-values corresponding to $\mathbf{e}_k$.

Assume that the vector $\mathbf{x}$ has coordinates ξ_k in the basis $\{\mathbf{e}_k\}$:

$$\mathbf{x} = \sum_{k=1}^n \xi_k \mathbf{e}_k. \quad (7.48)$$

Then, since $\mathbf{e}_k$ are the eigenvectors of the operator $\mathbf{A}$, we evidently have

$$A\mathbf{x} = \sum_{k=1}^n \lambda_k \xi_k \mathbf{e}_k. \quad (7.49)$$

Since the basis $\{\mathbf{e}_k\}$ is orthonormal, we obtain the following expression for the scalar product $(A\mathbf{x}, \mathbf{x})$ from relations (7.48) and (7.49):

$$(A\mathbf{x}, \mathbf{x}) = \sum_{k=1}^n \lambda_k \xi_k^2. \quad (7.50)$$

From this and from relation (7.47) we get (7.46). We have proved the theorem.

7.6.3. Simultaneous reduction of two quadratic forms to the sum of squares in a linear space.

We shall now prove an important theorem on simultaneous reduction of two quadratic forms to the sum of squares in an arbitrary (not necessarily Euclidean) real linear space.

Theorem 7.9. *Assume that $A(\mathbf{x}, \mathbf{y})$ and $B(\mathbf{x}, \mathbf{y})$ are symmetric bilinear forms defined in the real linear space V . Assume furthermore that the inequality $B(\mathbf{x}, \mathbf{x}) > 0$ is valid for all $\mathbf{x} \in V, \mathbf{x} \neq 0$, (i.e. the quadratic form $B(\mathbf{x}, \mathbf{x})$ is positive definite). Then in the space V there is a basis $\{\mathbf{e}_k\}$ such that the quadratic forms $A(\mathbf{x}, \mathbf{x})$ and $B(\mathbf{x}, \mathbf{x})$ can be represented as*

$$A(\mathbf{x}, \mathbf{x}) = \sum_{k=1}^n \lambda_k \xi_k^2, \quad (7.51)$$

$$B(\mathbf{x}, \mathbf{x}) = \sum_{k=1}^n \xi_k^2, \quad (7.52)$$

where ξ_k are the coordinates of the vector $\mathbf{x}$ in the basis $\{\mathbf{e}_k\}$.

Proof. In accordance with the remark made at the end of 7.2, in a finite-dimensional real space scalar product can be defined by means of the bilinear form $B(\mathbf{x}, \mathbf{y})$ which is the polar of the positive definite quadratic form $B(\mathbf{x}, \mathbf{x})$.

We can, therefore, introduce a scalar product $(\mathbf{x}, \mathbf{y})$ of the vectors $\mathbf{x}$ and $\mathbf{y}$ in the linear space V setting

$$(\mathbf{x}, \mathbf{y}) = B(\mathbf{x}, \mathbf{y}) \quad (7.53)$$

Thus, V is a Euclidean space with the scalar product (7.53). By Theorem 7.11, we can indicate an orthonormal basis $\{\mathbf{e}_k\}$ and real numbers λ_k such that the quadratic form $A(\mathbf{x}, \mathbf{x})$ is in the form (7.51) in that basis.

On the other hand, in any orthonormal basis, the scalar product $(\mathbf{x}, \mathbf{x})$ which is equal, according to (7.53), to $B(\mathbf{x}, \mathbf{x})$, can be represented as the sum of the squares of the coordinates of the vector $\mathbf{x}$. Thus, we have also substantiated the representation $B(\mathbf{x}, \mathbf{x})$ in the form (7.52) in the basis $\{\mathbf{e}_k\}$. And this is what we wished to prove.

Remark. It follows directly from the theorem we have proved that any quadratic form in an arbitrary real linear space can be reduced to a canonical form. However, the method of this reduction is, in general, more complicated than the methods presented in 7.3 since it requires finding all the eigenvectors of a self-adjoint operator (see Chapter 6).

7.6.4. Extremal properties of a quadratic form. Let us consider an arbitrary differentiable function f defined on a certain smooth surface S (see the definition of a smooth surface in Chapter 5 of our *Fundamentals of Mathematical Analysis*, Part 2). We say that the point $\mathbf{x}_0$ of the surface S is a stationary point of the function f if at the point $\mathbf{x}_0$ the derivative of the function f in any direction on the surface S is equal to zero. In particular, the points of extremum of the function f are its stationary points.

The value $f(\mathbf{x}_0)$ of the function f at the stationary point $\mathbf{x}_0$ is called its **stationary value**. The stationary point $\mathbf{x}_0$ of the function f is sometimes called its **critical point**, and the value, $f(\mathbf{x}_0)$ its **critical value**. We shall discuss here stationary and, in particular, extremal values of the quadratic form $B(\mathbf{x}, \mathbf{x})$ on a sphere of unit radius in the Euclidean space V and show how these values are connected with the eigenvalues of the self-adjoint operator A by means of which the symmetric bilinear form $B(\mathbf{x}, \mathbf{y})$, which is the polar of the quadratic form $B(\mathbf{x}, \mathbf{x})$, can be represented as

$$B(\mathbf{x}, \mathbf{y}) = (A\mathbf{x}, \mathbf{y}). \quad (7.54)$$

A unit sphere in V is the set of vectors $\mathbf{x} \in V$ which satisfy the equation

$$(\mathbf{x}, \mathbf{x}) = 1 \quad \text{or} \quad \|\mathbf{x}\| = 1. \quad (7.55)$$

To simplify the arguments, we shall use the conclusions made in 7.6.3 concerning the reduction of a quadratic form to the sum of squares.

Thus, suppose $B(\mathbf{x}, \mathbf{x})$ is a quadratic form, $B(\mathbf{x}, \mathbf{y})$ is a bilinear forms which is the polar of that form and A is a self-adjoint operator related to $B(\mathbf{x}, \mathbf{y})$ as (7.54).

By Theorem 7.8, in the orthonormal basis $\{\mathbf{e}_k\}$ consisting of the eigenvectors of the operator A the quadratic form $B(\mathbf{x}, \mathbf{x})$ has the form

$$B(\mathbf{x}, \mathbf{x}) = \sum_{k=1}^n \lambda_k \xi_k^2, \quad (7.56)$$

where ξ_k are the coordinates of the vector $\mathbf{x}$ in the basis $\{\mathbf{e}_k\}$, and λ_k are the eigenvalues of the operator A . We agree to enumerate these eigenvalues in a decreasing order:

$$\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n. \quad (7.57)$$

Note that in the chosen basis the unit sphere, defined by equation (7.55), in the coordinates of the vector $\mathbf{x}$ is specified by the equation

$$\sum_{k=1}^n \xi_k^2 - 1 = 0. \quad (7.58)$$

Let us prove the following theorem,

Theorem 7.10. *The stationary values of the quadratic form $B(\mathbf{x}, \mathbf{x})$ on unit sphere (7.55) are equal to the eigenvalues λ_k of the operator A . These stationary values can be attained, in particular, on the unit eigenvectors $\mathbf{e}_k$ of the operator A .*

Proof. Since we speak of the stationary values of the function $B(\mathbf{x}, \mathbf{x})$ under the condition $(\mathbf{x}, \mathbf{x}) = 1$, i.e. of the conditional extremum of that function, we can apply Lagrange's method of undetermined multipliers (see our *Fundamentals of Mathematical Analysis*, Part 1, 15.5.2). Let us form Lagrange's function

$\Psi(\xi_1, \xi_2, \dots, \xi_n)$ for the function $B(\mathbf{x}, \mathbf{x})$ using its expression (7.56) in the given basis $\{\mathbf{e}_k\}$ and taking into account that the connection relation has the form. (7.58). We get

$$\Psi = \sum_{k=1}^n \lambda_k \xi_k^2 - \lambda \left(\sum_{k=1}^n \xi_k^2 - 1 \right) \quad (7.59)$$

where λ is Lagrange's undetermined multiplier.

Recall that if we choose λ in (7.59) such that under condition; (7.58) the relations

$$\frac{\partial \Psi}{\partial \xi_k} = 0, k = 1, 2, \dots, n \quad (7.60)$$

hold true, then at the points of the sphere (7.58), which correspond to those values of λ , the function $B(\mathbf{x}, \mathbf{x})$ (the quadratic form $B(\mathbf{x}, \mathbf{x})$) has a stationary value.

Thus the problem concerning the stationary values of $B(\mathbf{x}, \mathbf{x})$ on the sphere $(\mathbf{x}, \mathbf{x}) = 1$ reduces to the investigation of the behaviour of the system of equations (7.58), (7.60) with respect to the variables λ and the coordinates $\xi_1, \xi_2, \dots, \xi_n$ of the vector $\mathbf{x}$. Note that in this case $\xi_1, \xi_2, \dots, \xi_n$ are the coordinates of the vector $\mathbf{x}$ on which $B(\mathbf{x}, \mathbf{x})$ has a stationary value.

Since $\frac{\partial \Psi}{\partial \xi_k} = 2(\lambda_k - \lambda) \xi_k$, system (7.58), (7.60) in question assumes the form

$$\left. \begin{aligned} \sum_{k=1}^n \xi_k^2 &= 1 \\ (\lambda_k - \lambda) \xi_k &= 0, k = 1, 2, \dots, n \end{aligned} \right\} \quad (7.61)$$

Let system (7.61) have a solution

$$\lambda = \tilde{\lambda}, \quad \tilde{\mathbf{x}} = (\tilde{\xi}_1, \tilde{\xi}_2, \dots, \tilde{\xi}_n).$$

Multiplying each relation $(\lambda_k - \tilde{\lambda})\tilde{\xi}_k = 0$ by $\tilde{\xi}_k$, we then add the resulting relations together and, taking into account that $\sum_{k=1}^n \tilde{\xi}_k^2 = 1$, we get, according to

(7.56) the following value for $\tilde{\lambda}$:

$$\tilde{\lambda} = \sum_{k=1}^n \lambda_k \tilde{\xi}_k^2 = B(\tilde{\mathbf{x}}, \tilde{\mathbf{x}}).$$

Thus if $\tilde{\lambda}$ and $\tilde{\mathbf{x}} = (\tilde{\xi}_1, \tilde{\xi}_2, \dots, \tilde{\xi}_n)$ is a solution of system (7.61), then $\tilde{\lambda}$ is equal to the value of the quadratic form $B(\mathbf{x}, \mathbf{x})$ on the vector $\tilde{\mathbf{x}} = (\tilde{\xi}_1, \tilde{\xi}_2, \dots, \tilde{\xi}_n)$ on which this form has a stationary value.

It is easy to see that the following values of the unknowns λ and ξ_i are solutions of system (7.61):

$$\lambda = \lambda_k, \xi_1 = 0, \dots, \xi_{k-1} = 0, \xi_k = 1, \xi_{k+1} = 0, \dots, \xi_n = 0,$$

$$k = 1, 2, \dots, n.$$

These solutions are, evidently, the eigenvalues λ_k and the coordinates of the corresponding eigenvectors $\mathbf{e}_k$. We have proved the theorem.

Remark. We have just found out that the eigenvalues λ_k are the stationary values of the quadratic form $B(\mathbf{x}, \mathbf{x})$ on the sphere $(\mathbf{x}, \mathbf{x}) = 1$.

It turns out that the numbers λ_1 and λ_n (under the condition (7.57)) are the greatest and the least value of $B(\mathbf{x}, \mathbf{x})$, respectively, on the sphere $(\mathbf{x}, \mathbf{x}) = 1$ (we have established above that these values can be attained).

To verify the validity of the remark, it is sufficient to replace all λ_k in (7.56) first by λ_n and then by λ_1 and use relations (7.57) and (7.58). Then we evidently obtain inequalities $\lambda_n \leq B(\mathbf{x}, \mathbf{x}) \leq \lambda_1$.

7.7. Second-Order Hypersurface

In this section we shall get acquainted with the main types of second-order hypersurfaces and indicate the means for investigating those surfaces.

7.7.1. The concept of a second-order hypersurface. Let V be an n -dimensional real Euclidean space.

For the sake of geometrical vividness, we shall call the vectors $\mathbf{x}$ of that space its **points**.

*A second-order **hypersurface** (hyperquadric) S is a locus of points $\mathbf{x}$ satisfying an equation of the form*

$$A(\mathbf{x}, \mathbf{x}) + 2B(\mathbf{x}) + c = 0, \quad (7.62)$$

where $A(\mathbf{x}, \mathbf{x})$ ($\mathbf{x}, \mathbf{x}$) is a quadratic form which is not identically zero, $B(\mathbf{x})$ is a linear form and c is a real number.

Equation (7.62) is a **general equation of a second-order hypersurface**.

We isolate an orthonormal basis $\{\mathbf{e}_k\}$ in the space V , and designate the coordinates of the vector $\mathbf{x}$ (point $\mathbf{x}$) in that basis as $((x_1, x_2, \dots, x_n))$. Then (see 7.1.2) the quadratic form $A(\mathbf{x}, \mathbf{x})$ can be represented as

$$A(\mathbf{x}, \mathbf{x}) = \sum_{j, k=1}^n a_{jk} x_j x_k, \quad (7.63)$$

where

$$a_{jk} = A(\mathbf{e}_j, \mathbf{e}_k), \quad (7.64)$$

and $A(\mathbf{e}_j, \mathbf{e}_k)$ is the value of the symmetric bilinear form $A(\mathbf{x}, \mathbf{y})$ which is the polar of the quadratic form $A(\mathbf{x}, \mathbf{x})$, on the vectors $\mathbf{e}_j$ and $\mathbf{e}_k$.

In the indicated basis $\{\mathbf{e}_k\}$ the linear form $B(\mathbf{x})$ can be represented in the form

$$B(\mathbf{x}) = \sum_{k=1}^n b_k x_k. \quad (7.65)$$

Thus, the general equation of a second-order hypersurface in the Euclidean space V with the isolated basis $\{\mathbf{e}_k\}$ can be represented in the form

$$\sum_{j,k=1}^n a_{jk} x_j x_k + 2 \sum_{j,k=1}^n b_k x_k + c = 0. \quad (7.66)$$

Let us agree on the following terminology.

We shall call the term $A(\mathbf{x}, \mathbf{x}) = \sum_{j,k=1}^n a_{jk} x_j x_k$ the **group of the leading terms** of equation (7.62) or (7.66).

The group of terms $B(\mathbf{x}) + c = \sum_{k=1}^n b_k x_k + c$ will be called the **linear part** of equation (7.62) or (7.66).

In what follows we shall consider the matrices

$$A = \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \dots & \dots & \dots \\ a_{n1} & \dots & a_{nn} \end{pmatrix} \text{ and } B = \begin{pmatrix} a_{11} & \dots & a_{1n} & b_1 \\ \dots & \dots & \dots & \dots \\ a_{n1} & \dots & a_{nn} & b_n \\ b_1 & \dots & b_n & c \end{pmatrix} \quad (7.67)$$

and the determinants $\det A$ and $\det B$ of those matrices.

The investigation of second-order hypersurfaces will be carried out by the method similar to that used in analytic geometry in investigating curves and surfaces of second order defined by general equations.

The method consists in the fact that by choosing a special Cartesian system of coordinates on a plane (for second-order curves) or in space (for second-order surfaces) the equation of a curve or a surface is simplified as far as possible. Then, in the process of the investigation of that equation the geometrical properties of the curve or the surface are elucidated. In addition, the enumeration of all possible types of the elementary equations of second-order curves and surfaces enables us to classify them.

To use this method in a multi-dimensional case, we must first investigate the transformations (mappings) of an n -dimensional Euclidean space which are analogs of the transformations of rectangular Cartesian coordinates in the case of two or three dimensions.

In an n -dimensional case, transformations of this kind are translations and the transformations of bases which map an orthonormal basis into a new orthonormal basis. Exact definitions of these transformations will be given in 7.7.2.

A second-order hypersurface treated as a geometrical object of the space V does not, evidently, change when transformations of the kind indicated above are performed. We shall later ascertain that for every equation of the form (7.62) (or (7.66)), we can choose an origin and an orthonormal basis in V such that this equation written in the coordinates with respect to the new basis will be very simple and, therefore, as in the case of two or three dimensions, we can indicate the geometrical characteristics of such surfaces and classify them.

7.77.2. Translations in a Euclidean space. Transformations of orthonormal bases into orthonormal bases. A translation or parallel displacement in the Euclidean space V is a transformation defined by the relation

$$\mathbf{x} = \mathbf{x}' + \mathring{\mathbf{x}}, \quad (7.68)$$

where $\mathring{\mathbf{x}}$ is a fixed point called the origin of coordinates, or simply origin.

Let the points $\mathbf{x}$, $\mathbf{x}'$ and $\mathring{\mathbf{x}}$ possess co-ordinates equal respectively to $(x_1, x_2, \dots, x_n)$, $(x'_1, x'_2, \dots, x'_n)$, and $(\mathring{x}_1, \mathring{x}_2, \dots, \mathring{x}_n)$.

Then, in coordinates, the translation is defined by the formulas

$$x_k = x'_k + \mathring{x}_k, \quad k = 1, 2, \dots, n. \quad (7.69)$$

Note that a translation does not change any fixed basis.

Let us now elucidate the characteristics of the transformation of an orthonormal basis into an orthonormal basis.

We assume that the orthonormal basis $\{\mathbf{e}_k\}$ is transformed into a new orthonormal basis $\{\mathbf{e}'_k\}$. Let us express every vector $\mathbf{e}'_k$ in terms of the vectors $\{\mathbf{e}_k\}$. We obtain

(7.70)

We designate the matrix of transformation (7.70) as P :

(7.71)

Since the basis $\{e_k\}$ and $\{e'_k\}$ are orthonormal, we get from (7.70), by means of at scalar multiplication of e'_j elt by e'_k ,

(7.72)

Let us now consider the transposed matrix P' , i.e. the matrix obtained from P by interchanging rows and columns.

It is evident, according to (7.72), that

(7.73)

where I is an identity matrix.

Equations (7.73) show that the matrix P' is an inverse of P , i.e.

(1.74)

Let us now assume that we are considering a transformation of the orthonormal basis $\{\mathbf{e}_k\}$ by formulas (7.70) and that the matrix P of that transformation satisfies condition (7.73) (or, which is the same, (7.74)).

$$B'(\mathbf{x}') = A(\mathbf{x}', \mathring{\mathbf{x}}) B(\mathbf{x}'), \quad (7.77)$$

$$c' = A(\mathring{\mathbf{x}}, \mathring{\mathbf{x}}) + 2B(\mathring{\mathbf{x}}) + c. \quad (7.78)$$

* The quadratic form $A(\mathbf{x}, \mathbf{x})$ is connected with the symmetric bilinear form $A(\mathbf{x}, \mathbf{y})$ which is a polar of the form $A(\mathbf{x}, \mathbf{x})$. The bilinear form $A(\mathbf{x}, \mathbf{y})$ is linear with respect to the arguments $\mathbf{x}$ and $\mathbf{y}$. The expression $A(\mathbf{x}', \mathring{\mathbf{x}})$ appearing in what follows is the value of the form $A(\mathbf{x}, \mathbf{y})$ on the vectors $\mathbf{x}'$ and $\mathring{\mathbf{x}}$.

Let us write the resulting formulas in coordinates.

Suppose the coordinates of the points $\mathbf{x}'$ and $\mathring{\mathbf{x}}$ are equal to $x'_1, x'_2, \dots, x'_n$ and $\mathring{x}_1, \mathring{x}_2, \dots, \mathring{x}_n$ respectively. Since the basis $\{\mathbf{e}_k\}$ does not change in translation, the quadratic form $A(\mathbf{x}', \mathbf{x}')$ can be written as follows:

$$A(\mathbf{x}', \mathbf{x}') = \sum_{j, k=1}^n a_{jk} x'_j x'_k \quad (7.79)$$

(note that the coefficients $a_{jk} = A(\mathbf{e}_j, \mathbf{e}_k)$ do not change since the base vectors $\mathbf{e}_k$ remain unchanged).

Consequently, we can make an important conclusion: *in translation the group of the leading terms retains its form.*

Let us now consider formulas (7.77) and (7.78). Since

$$A(\mathbf{x}', \mathring{\mathbf{x}}) = \sum_{k=1}^n \left(\sum_{j=1}^n a_{jk} \mathring{x}_j \right) x'_k,$$

$$B(\mathbf{x}') = \sum_{k=1}^n b_k x'_k,$$

$$A(\mathring{\mathbf{x}}, \mathring{\mathbf{x}}) = \sum_{j, k=1}^n a_{jk} \mathring{x}_j \mathring{x}_k,$$

$$B(\mathring{\mathbf{x}}) = \sum_{k=1}^n b_k \mathring{x}_k,$$

formula (7.77), assumes the form

$$B'(\mathbf{x}') = \sum_{k=1}^n b'_k x'_k = \sum_{k=1}^n \left[\sum_{j=1}^n a_{jk} \dot{x}_j + b_k \right] x'_k, \quad (7.80)$$

and formula (7.78) is written as

$$c' = \sum_{j,k=1}^n a_{jk} \dot{x}_j \dot{x}_k + 2 \sum_{j,k=1}^n b_k \dot{x}_k + c. \quad (7.81)$$

Thus, equation (7.76) in coordinates has the form

$$\sum_{j,k=1}^n a_{jk} x'_j x'_k + 2 \sum_{k=1}^n b'_k x'_k + c' = 0 \quad (7.82)$$

We shall need an expression for c' different from (7.81). Let us write (7.81) as follows:

$$c' = \sum_{k=1}^n \left[\sum_{j=1}^n a_{jk} \dot{x}_j + b_k \right] \dot{x}_k + \sum_{k=1}^n b_k \dot{x}_k + c. \quad (7.82)$$

Keeping in mind that the coefficients b'_k are expressed, as follows from (7.80), by the formulas

$$b'_k = \sum_{j=1}^n a_{jk} \dot{x}_j + b_k \quad (7.84)$$

we get from (7.83) the required expression for c' :

$$c' = \sum_{k=1}^n (b'_k + b_k) \dot{x}_k + c. \quad (7.85)$$

7.7.4. Transformation of a general equation for a second-order hypersurface when passing from an orthonormal basis to an orthonormal one. Suppose the orthonormal basis $\{\mathbf{e}_k\}$ is transformed into a new orthonormal basis $\{\mathbf{e}'_k\}$ by formulas (7.70), and P is an orthogonal matrix of that transformation (see (7.71)). Then, in accordance with Remark 2 of this section, the coordinates x_k and x'_k of a point in the

bases $\{e_k\}$ and $\{e'_k\}$ are related as (7.75). Substituting the expression for x_k from (7.75) into the left-hand side of equation (7.66) and taking into account that because of the homogeneity of relations (7.75) the group of the leading terms and the linear part of equation (7.66) are transformed autonomously we get the following expression for the general equation of a second-order hypersurface in the coordinates x'_k of the points in the transformed basis $\{e'_k\}$:

$$\sum_{j,k=1}^n a'_{jk} x'_j x'_k + 2 \sum_{k=1}^n b'_k x'_k + c' = 0. \quad (7.86)$$

In accord with the autonomy of transformation of the group of the leading terms, there hold equalities

$$\left. \begin{aligned} \sum_{j,k=1}^n a'_{jk} x'_j x'_k &= \sum_{j,k=1}^n a_{jk} x_j x_k \\ \sum_{k=1}^n b'_k x'_k &= \sum_{k=1}^n b_k x_k, \\ c' &= c. \end{aligned} \right\} \quad (7.87)$$

Directing our attention to the first formula (7.87), we see that to determine the coefficients a'_{jk} we can use the rule for the transformation of the coefficients of quadratic form in passing to a new basis. Namely, if we use the designation A' for the matrix of the quadratic form $A(x, x)$ in the basis $\{e'_k\}$, then, according to Theorem 7.2 and the relation $P' = P^{-1}$, we get the following relation between the matrices A and A' of the form $A(x, x)$ in the bases $\{e_k\}$ and $\{e'_k\}$:

$$A' = P^{-1} A P \quad (7.88)$$

(recall that P is a matrix of an orthogonal transformation).

We shall now regard the matrix A' as the matrix of a certain linear operator A in the basis $\{e'_k\}$, and the matrix P^{-1} as the matrix of the passage from the basis $\{e'_k\}$ to

the basis $\{\mathbf{e}_k\}$. Then, according to Theorem 5.7 (see 5.2.2), the matrix A can be considered to be the matrix of the linear operator A in the basis $\{\mathbf{e}_k\}$.

In other words, *when we transform an orthonormal basis into an orthonormal one, the matrix of a quadratic form changes as the matrix of a certain linear operator.*

This inference will be used in 7.7.5.

Remark. Note that the operator A whose matrix in an orthonormal basis coincides with the matrix of the quadratic form $A(\mathbf{x}, \mathbf{x})$ is self-adjoint.

To prove this statement, we present the following arguments. Let $A(\mathbf{x}, \mathbf{x})$ be a quadratic form and $A(\mathbf{x}, \mathbf{y})$, a symmetric bilinear form which is the polar of $A(\mathbf{x}, \mathbf{x})$. According to Theorem 7.8, the bilinear form $A(\mathbf{x}, \mathbf{y})$ can be represented as

$$A(\mathbf{x}, \mathbf{y}) = (A\mathbf{x}, \mathbf{y}),$$

where A is a *self-adjoint operator*.

The quadratic form $A(\mathbf{x}, \mathbf{x})$ can, therefore, be represented as

$$A(\mathbf{x}, \mathbf{x}) = (A\mathbf{x}, \mathbf{x}).$$

Let us prove that in the orthonormal basis $\{\mathbf{e}_k\}$ the matrices of the operator A and of the quadratic form coincide. That will be the proof of the statement of the remark.

Suppose a_{jk} are the elements of the matrix of the form $A(\mathbf{x}, \mathbf{x})$ and $\tilde{a}_{jk}$ are the elements of the matrix of the operator A in the basis $\{\mathbf{e}_k\}$. According to 7.1.2, $\tilde{a}_{jk} = A(\mathbf{e}_j, \mathbf{e}_k)$ and, according to 5.2.1, formula (5.13), the elements a_{jk} can be found from the equations $A\mathbf{e}_j = \sum_{p=1}^n \tilde{a}_{jp} \mathbf{e}_p$.

We multiply scalarly both sides of the last relation by $\mathbf{e}_k$. Then, taking the orthonormality of the basis $\{\mathbf{e}_k\}$ into account, we get $(A\mathbf{e}_j, \mathbf{e}_k) = \tilde{a}_{jk}$. Since

$A(\mathbf{e}_j, \mathbf{e}_k) = (A\mathbf{e}_j, \mathbf{e}_k)$, we have $a_{jk} = \tilde{a}_{jk}$. We have proved the statement of the remark.

7.7.5. Invariants of the general equation of a second-order hypersurface. The invariant of general equation (7.62) (or (7.66)) of a second-order hypersurface relative to translations and transformations of the orthogonal bases into orthogonal ones is the function $f(a_{11}, a_{12}, \dots, a_{nn}, b_1, \dots, b_n, c)$ of the coefficients of this equation whose value does not change under the indicated transformations of a space.

Let us prove the following statement.

Theorem 7.11. *The invariants of general equation (7.62) (or (7.66)) of a second-order hypersurface are the coefficients of the characteristic polynomial of the matrix A of the quadratic form $A(\mathbf{x}, \mathbf{x})$ and the determinant $\det B$ of the matrix B in relation (7.67). In particular, the invariants are $\det A$ and the trace $a_{11} + a_{22} + \dots + a_{nn}$ of the matrix A .*

Proof. It is evidently sufficient to prove the invariance of the quantities enumerated in the hypothesis separately for the translation and transformation of an orthonormal basis into an orthonormal one.

Let us first consider the *translation*. We established in 7.3 that the group of the leading terms retained its form under this transformation (see formula (7.79)). Therefore, the matrix A does not change and, consequently, neither does the characteristic polynomial of that matrix.

Let us prove the invariance of $\det B$.

As a result of the translation (7.68) (or (7.69)), the matrix is transformed into the matrix B' , whose determinant, according to (7.82), has the form

$$\det B' = \begin{vmatrix} a_{11} & \dots & a_{1n} & b'_1 \\ \dots & \dots & \dots & \dots \\ a_{n1} & \dots & a_{nn} & b'_n \\ b'_1 & \dots & b'_n & c' \end{vmatrix}, \quad (7.89)$$

where b'_k and c' can be determined from (7.84) and (7.85).

Let us subtract from the last $(n + 1)$ th row of the determinant (7.89) the elements of the first row c multiplied by $\dot{x}_1$, then the elements of the second row multiplied by $\dot{x}_2$, and so on, and finally the elements of the n th row multiplied by $\dot{x}_n$. Since the determinant remains unchanged under these transformations, we can use (7.84) and (7.85) and obtain a relation

$$\det B' = \begin{vmatrix} a_{11} & \dots & a_{1n}b'_1 \\ \dots & \dots & \dots \\ a_{n1} & \dots & a_{nn}b'_n \\ b_1 & \dots & b'_n \left(\sum_{k=1}^n b_k \dot{x}_k + c \right) \end{vmatrix} \quad (7.90)$$

Let us now subtract from the elements of the last $(n+1)$ th column of determinant (7.90) the elements of the first column multiplied by $\dot{x}_1$, then the elements of the second column multiplied by $\dot{x}_2$, and so on, and, finally, the elements of the n th column multiplied by $\dot{x}_n$. Since the determinant does not change under these transformations, we can use the relation $a_{jk} = a_{kj}$ following from the symmetry of the form $A(x, y)$ and formula (7.84) and obtain $\det B$. Thus, we have proved the equality $\det B' = \det B$. Consequently, $\det B$ is invariant relative to translations.

Let us now consider the *transformation of an orthonormal basis into an orthonormal one*.

We shall first verify that the coefficients of the characteristic polynomial of the matrix A of the quadratic form are the invariants of the transformation being considered.

We established in 7.7.4 that when we pass to a new orthonormal basis, the matrix A changes as the matrix of a certain linear operator. But in that case, as follows from Remark 1 of 5.2.3, the coefficients of the characteristic polynomial of that matrix do not change in passing to a new basis.

In particular, the determinant $\det A$ and the trace $a_{11} + a_{22} + \dots + a_{nn}$ of the matrix A are invariants as the coefficients of a characteristic polynomial.

It remains to prove the invariance of the determinant $\det B$ in the transformation of an orthonormal basis into an orthonormal one. Let us prove this fact.

We apply the following technique. We introduce the designations $b_k = a_{k, n+1}$, $k = 1, 2, \dots, n$, $c = a_{n+1, n+1}$. Then equation (7.66) of a hypersurface can be written as follows:

$$\sum_{j, k=1}^{n+1} a_{jk} x_j x_k = 0, \quad (7.91)$$

where $x_{n+1} = 1$.

Let us consider the transformation of the variables $x_1, x_2, \dots, x_n, x_{n+1}$ into the variables $x'_1, x'_2, \dots, x'_n, x'_{n+1}$ under which the first n variables are transformed by formulas (7.75) and the variable x_{n+1} is transformed by the formula $x_{n+1} = x'_{n+1}$.

It is clear that this transformation of variables can be regarded as the transformation of coordinates in transforming the orthonormal basis $e_1, e_2, \dots, e_n, e_{n+1}$ of the $(n + 1)$ -dimensional Euclidean space, the matrix P of that transformation having the form

$$\tilde{P} = \begin{pmatrix} p_{11} & \dots & p_{n1} & 0 \\ \dots & \dots & \dots & \dots \\ p_{1n} & \dots & p_{nn} & 0 \\ 0 & \dots & 0 & 1 \end{pmatrix}. \quad (7.92)$$

It is easy to see that the matrix $\tilde{P}$ satisfies the condition $\tilde{P}' = \tilde{P}^{-1}$ and is, therefore, orthogonal. But then, according to 7.2.2, the orthonormal basis $e_1, e_2, \dots, e_n, e_{n+1}$ is transformed into an orthonormal basis by means of the matrix $\tilde{P}$. We have found

out above that under such transformation of the matrix B of the quadratic form, the determinant $\det B$ of that matrix is an invariant. We have proved the theorem.

Remark. It follows from the arguments presented in the proof of the theorem that the quantities $\text{rank } A$ and $\text{rank } B$ are also invariants of the general equation of a second-order hypersurface.

7.7.6. The centre of a second-order hypersurface. Let us try to find a translation under which general equation (7.76) would not contain the term $2B'(x')$ (or, if we turn to equation (7.82), the terms $2 \sum_{k=1}^n b'_k x'_k$).

To put in otherwise, we shall seek a translation (i.e. the coordinates $\overset{\circ}{x}_1, \overset{\circ}{x}_2, \dots, \overset{\circ}{x}_n$ of the point $\overset{\circ}{x}(\overset{\circ}{x})$ under which all the coefficients b'_k vanish. Directing our attention to formulas (7.84), we find that the required coordinates $\overset{\circ}{x}_1, \overset{\circ}{x}_2, \dots, \overset{\circ}{x}_n$ of the point $\overset{\circ}{x}$ are a solution of the following system of linear equations :

$$\sum_{j=1}^n a_{jk} \overset{\circ}{x}_j + b_k = 0, k = 1, 2, \dots, n. \quad (1.93)$$

Equations (7.93) are known as the **equations of the centre of a second-order hypersurface**, and the point $\overset{\circ}{x}$ with the coordinates $(\overset{\circ}{x}_1, \overset{\circ}{x}_2, \dots, \overset{\circ}{x}_n)$, where $(\overset{\circ}{x}_1, \overset{\circ}{x}_2, \dots, \overset{\circ}{x}_n)$ is a solution of system (7.93), is called the **centre** of that surface.

Here is an explanation of the term the “centre” of a second-order hypersurface. Assume that the origin is at the centre $\overset{\circ}{x}$, i.e. the required translation has been performed. Then the equation of the surface S assumes the form

$$\sum_{j,k=1}^n a_{jk} x'_j x'_k = 0. \quad (7.94)$$

Let the point x with the coordinates $(x'_1, x'_2, \dots, x'_n)$ lie on S . This means that its coordinates satisfy equation (7.94). The point x with the coordinates

$(-x'_1, -x'_2, \dots, -x'_n)$ which is symmetric with respect to the point $\mathbf{x}$ about the point $\mathring{\mathbf{x}}$, evidently also lies on S since its coordinates also satisfy equation (7.94).

Thus, if the hypersurface S has a centre, then the *points of S are located symmetrically pairwise about the centre.*

Remark 1. If the second-order hypersurface S has a centre, then the invariants $\det A$, $\det B$ and the constant term c' in equation. (7.94) are related as

$$\det B = c' \det A. \quad (7.95)$$

Indeed, we get for equation (7.94)

$$\det B = \begin{vmatrix} a_{11} & \dots & a_{1n} & 0 \\ \dots & \dots & \dots & \dots \\ a_{n1} & \dots & a_{nn} & 0 \\ 0 & \dots & 0 & c' \end{vmatrix}.$$

The last formula yields (7.95).

The existence of a centre I of a second-order hypersurface is connected with the solvability of the equations of centre (7.93).

*If the equations of a centre have a unique solution, then the hypersurface S is called a **central hypersurface**.*

Since the determinant of system (7.93) is equal to $\det A$, and the necessary and sufficient condition for a unique solution of that system to exist is that its determinant should differ from zero, we can infer that *for the hypersurface S to be central, it is necessary and sufficient that $\det A \neq 0$.*

Remark 2. If the origin is transferred to the centre of the central hypersurface S , then the equation of that hypersurface has the form

$$\sum_{j,k=1}^n a_{jk} x_j x_k + \frac{\det B}{\det A} = 0. \quad (7.96)$$

Indeed, with the origin transferred to the centre, the equation of the hypersurface assumes the form (7.94). Since $\det A \neq 0$ for the central hypersurface, we find

from (7.95) that $c' = \det B / \det A$. Substituting this expression for c' into (7.94), we obtain equation (7.96).p

7.7.7. The standard simplification of any equation of a second-order hypersurface by transforming an orthonormal basis. By Theorem 7.8, there is an orthonormal basis in which the quadratic form $A(x, x)$ can be written as the sum of squares. Let us designate this basis as $\{e'_k\}$ and the coordinates of the point x in that basis as $x'_1, x'_2, \dots, x'_n$. Besides that, we use the letters $\lambda_1, \lambda_2, \dots, \lambda_n$ to designate the eigenvalues of the self-adjoint operator A whose matrix in the orthonormal basis coincides with that of the quadratic form $A(x, x)$ (see the remark in 7.7.4).

Using now the conclusions of Theorem 7.8, we write the quadratic form $A(x, x)$ in the coordinates $x'_1, x'_2, \dots, x'_n$ of the point x in the basis $\{e'_k\}$ as follows:

$$A(x, x) = \sum_{p=1}^n \lambda_p^2 x_p'^2. \quad (7.97)$$

Let us pass from the basis $\{e_k\}$ to the basis $\{e'_k\}$. Since the formulas for the transformation of the coordinates of points are linear and homogeneous under such a transformation (see the remark in 7.7.2, formulas (7.75)), the group of the leading terms and the linear part of the equation of the hypersurface S are transformed autonomously. By virtue of this and equation (7.97), the equation of the hypersurface S in the basis $\{e'_k\}$ has the form:

$$\sum_{p=1}^n \lambda_p^2 x_p'^2 + 2 \sum_{k=1}^n b'_k x'_k + c = 0. \quad (7.98)$$

The reduction of any equation of the second-order hypersurface S to the form (7.98) is called the *standard simplification of that equation (by the transformation of the orthonormal basis)*.

7.7.8. Simplification of the equation of a central second-order hypersurface. Classification of central hypersurfaces. The conclusions made in 7.7.6 and 7.7.7 enable us to solve the problem of classifying all the central second-order

hypersurfaces. The solution of the problem will be carried out by the following scheme. First, by transferring the origin to the centre of the hypersurface (7.66), we reduce the equation to the form (7.96). Then we perform a standard simplification of equation (7.96). As a result, we evidently get according to (7.98) the following equation of the central second-order hyper-surface:

$$\lambda_1 x_1''^2 + \lambda_2 x_2''^2 + \dots + \lambda_n x_n''^2 + \frac{\det B}{\det A} = 0, \quad (7.99)$$

in which λ_k are the eigen numbers of the matrix A of the quadratic form $A(\mathbf{x}, \mathbf{x})$ in equation (7.62), and x'' are the coordinates of the point $\mathbf{x}$ in the final orthonormal basis $\{\mathbf{e}'_k\}$.

Note that *all the eigen numbers $\lambda_k, k = 1, 2, \dots, n$, are nonzero*.

Indeed, calculating $\det A$ for equation (7.99), we get $\det A = \lambda_1 \lambda_2 \dots \lambda_n$ and since $\det A \neq 0$ for a central surface, it evidently follows that $\lambda_k \neq 0$.

Next we agree to enumerate all positive eigen numbers of the matrix A by the first index and all negative eigen numbers by the subsequent indices. Thus, there is a number p such that

$$\begin{aligned} \lambda_1 > 0, \lambda_2 > 0, \dots, \lambda_p > 0, \\ \lambda_{p+1} < 0, \lambda_{p+2} < 0, \dots, \lambda_n < 0. \end{aligned}$$

We now introduce the following designations :

if $\operatorname{sgn} \frac{\det B}{\det A} \neq 0$, we set

$$\left. \begin{aligned} \left| \frac{\det B}{\det A} \right| \lambda_k &= \frac{1}{a_k^2} \text{ for } k = 1, 2, \dots, p, \\ \left| \frac{\det B}{\det A} \right| \lambda_k &= -\frac{1}{a_k^2} \text{ for } k = p+1, \dots, n, \end{aligned} \right\} \quad (7.100)$$

if $\operatorname{sgn} \frac{\det B}{\det A} = 0$, we set

$$\left. \begin{aligned} \lambda_k &= \frac{1}{a_k^2} \text{ for } k = 1, 2, \dots, p, \\ \lambda_k &= -\frac{1}{a_k^2} \text{ for } k = p+1, \dots, n. \end{aligned} \right\} \quad (7.101)$$

Then equation (7.99) can evidently be rewritten as follows (we replace the designation of the coordinates x_k'' by x_k) :

$$\frac{x_1^2}{a_1^2} + \frac{x_2^2}{a_2^2} + \dots + \frac{x_p^2}{a_p^2} - \frac{x_{p+1}^2}{a_{p+1}^2} - \dots - \frac{x_n^2}{a_n^2} + \operatorname{sgn} \frac{\det B}{\det A} = 0. \quad (7.102)$$

Equation (7.102) is known as the **canonical equation** of the central second-order hypersurface.

The quantities $a_k, k=1, 2, \dots, n$, are called the *semi-axes of the central second-order hypersurface*. They can be calculated by formulas (7.100) and (7.101).

We use equation (7.102) to classify central hypersurfaces as follows.

1°. $p = n, \operatorname{sgn} \frac{\det B}{\det A} = -1$. In this case, the hyper surface S is called an **(n – 1)-dimensional ellipsoid**.

The canonical equation of such an ellipsoid is usually written as

$$\frac{x_1^2}{a_1^2} + \dots + \frac{x_n^2}{a_n^2} = 1. \quad (7.103)$$

If $a_1 = a_2 = \dots = a_n = R$, then an $(n – 1)$ -dimensional ellipsoid is a sphere of radius R in an n -dimensional space.

Remark 1. In the case $p = 0, \operatorname{sgn} \frac{\det B}{\det A} = 1$, we also get an $(n – 1)$ -dimensional ellipsoid. It is evident that in this case equation (7.102) can be written in the form (7.103).

2°. $p = n$, $\operatorname{sgn} \frac{\det B}{\det A} = 1$. The hypersurface is imaginary and is called an **imaginary ellipsoid**.

Remark 2. It is evident that in the case $p = 0$, $\operatorname{sgn} \frac{\det B}{\det A} = -1$ we also obtain an imaginary ellipsoid.

3°. $0 < p < n$, $\operatorname{sgn} \frac{\det B}{\det A} \neq 0$. In this case the central hypersurfaces are called **hyperboloids**.

The geometrical characteristics of a hyperboloid depend on the ratios of the numbers p and n and the value of $\operatorname{sgn} \frac{\det B}{\det A}$.

4°. $\operatorname{sgn} \frac{\det B}{\det A} = 0$. In this case the central hypersurfaces are called **singular**.

Among singular hypersurfaces, note the so-called **singular ellipsoid** corresponding to the values $p = 0$ and $p = n$.

7.7.9. Simplification of the equation of a noncentral second-order hypersurface. Classification of noncentral hypersurfaces. Suppose that the hypersurface S , specified by equation (7.62) is not central, i.e.

$$\det A = 0. \quad (7.104)$$

We carry out a standard simplification of equation (7.62). As a result, the equation assumes the form (7.98). We calculate $\det A$ using (7.98) (this is possible since $\det A$ is an invariant). Taking (7.104) into account, we obtain

$$\det A = \lambda_1 \lambda_2 \dots \lambda_n = 0.$$

Thus, at least one eigenvalue λ_k of the matrix A is zero. It should be emphasized that not all the eigenvalues are zero, since otherwise the quadratic form $A(x, x)$ would be identically equal to zero, whereas we assume (see 7.1.1) that the form is nonzero.

We leave in expression (7.98) only those terms in the first sum that correspond to the nonzero eigenvalues and then renumber the base vectors so that all the nonzero

eigenvalues $\lambda_1, \lambda_2, \dots, \lambda_n$ correspond to the first p base vectors $e'_1, \dots, e'_p$ (note that $p = \text{rank } A$). After that, equation (7.98) can evidently be rewritten as

$$\sum_{k=1}^n \lambda_k x_k'^2 + 2 \sum_{k=1}^n b'_k x'_k + 2 \sum_{k=p+1}^n b'_k x'_k + c = 0 \quad (7.105)$$

(here $0 < p < n$, $\lambda_1 \neq 0, \dots, \lambda_p \neq 0$; in addition, we especially isolated the first p terms of the second sum in equation (7.98)).

Let us now perform the following transformations.

1°. For every number $k, 1 \leq k \leq p$, we unite the terms with that number from the first and the second sum in (7.105) and then make the following transformations (taking into account that $\lambda_k \neq 0$):

$$\begin{aligned} \lambda_k x_k'^2 + 2b'_k x'_k &= \lambda_k \left(x_k'^2 + 2 \frac{b'_k}{\lambda_k} x'_k + \frac{b_k'^2}{\lambda_k^2} \right) \\ &\quad - \frac{b_k'^2}{\lambda_k} = \lambda_k \left(x'_k + \frac{b'_k}{\lambda_k} \right)^2 - \frac{b_k'^2}{\lambda_k}. \end{aligned}$$

After these transformations, (7.105) will evidently look like

$$\sum_{k=1}^n \lambda_k \left(x'_k + \frac{b'_k}{\lambda_k} \right)^2 + \sum_{k=p+1}^n b'_k x'_k + c' = 0, \quad (7.106)$$

where the constant c' is defined by the equation

$$c' = c - \sum_{k=1}^n \frac{b_k'^2}{\lambda_k}. \quad (7.107)$$

Next we perform a translation by the formulas

$$x_k'' = x'_k + \frac{b'_k}{\lambda_k}, \quad k = 1, 2, \dots, p; \quad x_k'' = x'_k, \quad k = p+1, \dots, n.$$

As a result, equation (7.106) passes into an equation

$$\sum_{k=1}^p \lambda_k x_k'' + \sum_{k=p+1}^n b_k' x_k'' + c' = 0, \quad (7.108)$$

where c' is defined by formula (7.107).

2°. We shall now seek a transformation of the orthonormal basis $\{\mathbf{e}_k'\}$ under which the first p base vectors $\mathbf{e}_1', \dots, \mathbf{e}_p'$ do not change. At the expense of a change in the

base vectors $\mathbf{e}_{p+1}', \dots, \mathbf{e}_n'$ we shall try to reduce the term $2 \sum_{k=p+1}^n b_k' x_k''$ to the

form $2\mu x_n'''$, where x_n''' is the n th coordinate in the new basis. Note that when we perform such a transformation, the constant term c' does not change.

It should be pointed out first that if all the coefficients b_k' in (7.108) are zero, then the aim of the transformation of 2° is attained, i.e. the term $2 \sum_{k=p+1}^n b_k' x_k''$ has the

form $2\mu x_n''$, where $\mu = 0$.

Thus we assume that at least one of the coefficients b_k' in the sum $\sum_{k=p+1}^n b_k' x_k''$ is

nonzero. Then we can regard that sum as a certain linear form $B''(x)$ defined in the subspace V'' which is a linear closure of the vectors $\mathbf{e}_{p+1}', \dots, \mathbf{e}_n'$. According

to the lemma of 5.4.1, this form can be represented in the indicated subspace as $B''(x) = (\mathbf{h}, x)$, where $\mathbf{h}$ is some vector of the subspace V'' . If now, in the

subspace V'' , we direct the unit vector $\mathbf{e}_n''$ along the vector $\mathbf{h}$ so that $\mathbf{h} = \mu \mathbf{e}_n''$, and choose the vectors $\mathbf{e}_{p+1}'', \dots, \mathbf{e}_{n-1}''$ such that the system $\mathbf{e}_{p+1}'', \dots, \mathbf{e}_{n-1}'', \mathbf{e}_n''$ is a

basis in V'' then, evidently, $B''(x) = (\mathbf{h}, x) = \mu(\mathbf{e}_n'', x) = \mu x_n'''$ in that basis since $(\mathbf{e}_n'', x) = x_n'''$. Thus, choosing a basis in V'' in a way described above, we reduce

$\sum_{k=p+1}^n b_k' x_k''$ to the form $\mu x_n'''$.

Thus, there is a transformation of the basis $\mathbf{e}'_1, \dots, \mathbf{e}'_n$ into the orthonormal basis $\mathbf{e}''_1, \dots, \mathbf{e}''_n$ (under which the vectors $\mathbf{e}'_1, \dots, \mathbf{e}'_p$ remain unchanged) such that equation (7.108) assumes the form (we replace the designation of the coordinates x''' by x_k).

$$\sum_{k=1}^p \lambda_k x_k^2 + 2\mu x_n + c' = 0. \quad (7.109)$$

Note that the case $\mu = 0$ is not excluded in equation (7.109).

Equation (7.109) is called the **canonical equation of a noncentral second-order hypersurface**.

We can use equation (7.109) to **classify** noncentral hypersurfaces. Three cases are possible.

1°. $\mu \neq 0, p = \text{rank } A = n - 1$.

In this case, we write the last two terms in equation (7.109) in the form

$2\mu x_n + c' = 2\mu \left(x_n + \frac{c'}{2\mu} \right)$ and perform a translation along the axis x_n through the value $\frac{-c'}{2\mu}$. In order not to make the notation too cumbersome, we shall not change

the designation of the coordinates in this case. As a result, the canonical equation (7.109) assumes the form

$$\lambda_1 x_1^2 + \dots + \lambda_{n-1} x_{n-1}^2 + 2\mu x_n = 0. \quad (7.110)$$

The second-order hypersurfaces whose canonical equation has the form (7.110) are called **paraboloids**.

2°. $\mu = 0, p = \text{rank } A < n$.

In this case, canonical equation (7.109) can be rewritten as

$$\lambda_1 x_1^2 + \dots + \lambda_p x_p^2 + c' = 0. \quad (7.111)$$

In the subspace which is a linear closure of the vectors $e_1''', \dots, e_p'$ equation (7.111) is, evidently, the canonical equation of the central second-order surface S' . To get an idea of the hypersurface S in the whole space, we must place a plane, which is parallel to the plane V'' (a linear closure of the vectors $e_{p+1}', \dots, e_n'$) at each point of the surface S' . The locus of such planes forms a surface S . Thus, the surface S is a **central cylinder** with a directing surface S' defined by equation (7.111) and generating surfaces which are parallel to the plane V'' .

3°. $\mu \neq 0$, $p = \text{rank } A < n - 1$.

Acting as in 1°, we reduce the canonical equation (7.109) to the form

$$\lambda_1 x_1^2 + \dots + \lambda_p x_p^2 + 2\mu x_n = 0. \quad (1.112)$$

In the subspace which is a linear closure of the vectors $e_1', \dots, e_p', e_n'$, equation (7.112) evidently defines a paraboloid S' (see case 1°). To get an idea of the structure of the hypersurface in the whole space, we must place a plane which is parallel to the plane V'' (a linear closure of the vectors $e_{p+1}', \dots, e_{n-1}'$) at each point of S' . The locus of such planes forms a surface S . Thus, the surface S is a paraboloidal cylinder with a directing surface S' defined by equation (7.112) and generating planes which are parallel to the plane V'' .

Chapter – 8

TENSORS

In this chapter we consider important objects called tensors which are characterized, in every basis, by a collection of coordinates transformed, in a special way, in the passage from one basis to another. Tensors are widely employed in geometry, physics and mechanics. The concept of a tensor arises when we study various anisotropic phenomena (say, when we study the dependence of the distribution of the velocities of light propagation in a crystal on the direction of its propagation).

8.1. Transformation of Bases and Coordinates

This section dealing with the laws of transformation of coordinates in an arbitrary real Euclidean space E^n is of an auxiliary nature. The leading considerations arising in the process make the notion of a tensor introduced in the next section more visual.

8.1.1. Gram's determinants. We shall indicate here the means of elucidating the linear dependence of the system of vectors $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_k$ in a Euclidean space.

For that purpose, we introduce the so-called Gram's determinant, also called a Gramian, of the indicated system of vectors.

The Gramian of the system of vectors $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_k$ is the following determinant:

$$\begin{vmatrix} (\mathbf{e}_1, \mathbf{e}_1) & (\mathbf{e}_1, \mathbf{e}_2) & \dots & (\mathbf{e}_1, \mathbf{e}_k) \\ (\mathbf{e}_2, \mathbf{e}_1) & (\mathbf{e}_2, \mathbf{e}_2) & \dots & (\mathbf{e}_2, \mathbf{e}_k) \\ \dots & \dots & \dots & \dots \\ (\mathbf{e}_k, \mathbf{e}_1) & (\mathbf{e}_k, \mathbf{e}_2) & \dots & (\mathbf{e}_k, \mathbf{e}_k) \end{vmatrix}. \quad (8.1)$$

The following statement holds true.

Theorem 8.1. *For the system of vectors $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_k$ of the Euclidean space E^n to be linearly dependent, it is necessary and sufficient that Gramian (8.1) of that system be zero.*

Proof. (1) Necessity. Suppose the vectors $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_k$ are linearly dependent. Then one of them, say, $\mathbf{e}_k$, is a linear combination of the other vectors:

$$\mathbf{e}_k = \alpha_1 \mathbf{e}_1 + \alpha_2 \mathbf{e}_2 + \dots + \alpha_{k-1} \mathbf{e}_{k-1}.$$

Performing a scalar multiplication of the last relation by $\mathbf{e}_i, i = 1, 2, \dots, k$, we find that the last row of Gramian (8.1) is a linear combination of the first $k - 1$ rows. By Theorem 1.7, this determinant is zero. We have proved the necessity.

(2) Sufficiency. We assume that Gramian (8.1) is zero. Then its columns are linearly dependent, i.e. there are numbers $\beta_1, \beta_2, \dots, \beta_k$ which are not all equal to zero and such that the relations

$$\beta_1 (\mathbf{e}_i, \mathbf{e}_1) + \beta_2 (\mathbf{e}_i, \mathbf{e}_2) + \dots + \beta_k (\mathbf{e}_i, \mathbf{e}_k) = 0$$

are satisfied for $i = 1, 2, \dots, k$. Rewriting these relations in the form

$$(\mathbf{e}_i, \beta_1 \mathbf{e}_1 + \beta_2 \mathbf{e}_2 + \dots + \beta_k \mathbf{e}_k) = 0, \quad i = 1, 2, \dots, k,$$

we make sure that the vector $\beta_1 \mathbf{e}_1 + \beta_2 \mathbf{e}_2 + \dots + \beta_k \mathbf{e}_k$ is orthogonal to all the vectors $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_k$ i.e. is orthogonal to the linear closure L of those vectors. Since this vector belongs to L , it is equal to zero. Since not all the β_j are zero, this means that the vectors $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_k$ are linearly dependent. And this is what we set out to prove.

Corollary. *If the vectors $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_k$ are linearly independent, then the Gramian of those vectors is nonzero.*

Let us prove that in this case the Gramian is **positive**. Let L be a linear closure of the vectors $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_k$. It is evident that $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_k$ is a basis in L . Let us consider a bilinear symmetric form $A(\mathbf{x}, \mathbf{y})$ which is a scalar product $(\mathbf{x}, \mathbf{y}) : A(\mathbf{x}, \mathbf{y}) = (\mathbf{x}, \mathbf{y})$. The corresponding quadratic form $A(\mathbf{x}, \mathbf{x}) = (\mathbf{x}, \mathbf{x})$ is evidently a definite form and, therefore, in accordance with Theorem 7.6 (Sylvester's criterion), the determinant $\det(a_{ij})$ of its matrix (a_{ij}) in the basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_k$ is positive. But

this determinant is precisely Gramian (8.1) of the system $e_1, e_2, \dots, e_k$ since $a_{ij} = (e_i, e_j)$.

8.1.2. Reciprocal bases. Covariant and contravariant vector coordinates. Let $e_1, e_2, \dots, e_n$ be a basis in the Euclidean space E^n . The basis $e_1, e_2, \dots, e_n$ is said to be a **reciprocal** of the basis $e_i, i = 1, 2, \dots, n$, if the relations

$$(e_i, e^j) = \delta_i^j = \begin{cases} 1 & \text{for } i = j \\ 0 & \text{for } i \neq j \end{cases} \quad (8.2)$$

are satisfied for $i, j = 1, 2, \dots, n$.

The symbol δ_i^j is known as the **Kronecker delta**.

A question arises whether a reciprocal basis exists and is unique. The question is answered in the affirmative: *there is a unique reciprocal basis for any given basis $e_1, e_2, \dots, e_n$* . To prove this statement, we do the following. We assume that $x_1^j, x_2^j, \dots, x_n^j$ are the coordinates of the required vectors e^j in the basis e_i :

$$e^j = x_1^j e_1 + x_2^j e_2 + \dots + x_n^j e_n, \quad j = 1, 2, \dots, n. \quad (8.3)$$

Multiplying scalarly both sides of the last relations by e_i , we get, using (8.2),

$$x_1^j (e_i, e_1) + x_2^j (e_i, e_2) + \dots + x_n^j (e_i, e_n) = \delta_i^j \quad (8.4)$$

$$i, j = 1, 2, \dots, n.$$

For a fixed j , relations (8.4) can be regarded as a quadratic system of linear equations with respect to the unknown coordinates $x_1^j, x_2^j, \dots, x_n^j$ of the vector e^j in the basis e_i . Since the determinant of system (8.4) is the Gramian of the base vectors $e_1, e_2, \dots, e_n$, it follows, in accordance with the corollary of Theorem 8.1, that it is non-zero and, therefore, system (8.4) has a unique solution $x_1^j, x_2^j, \dots, x_n^j$ which is nonzero, since the system is non-homogeneous. Then we construct vectors e^j with the aid of relations (8.3) which evidently satisfy relations (8.2).

It remains to verify that the vectors $e^1, e^2, \dots, e^n$ form a basis.

Suppose a certain linear combination of these vectors is equal to

$$\text{zero: } \alpha_1 e^1 + \alpha_2 e^2 + \dots + \alpha_n e^n = 0.$$

Multiplying scalarly the last equation in succession by $e_1, e_2, \dots, e_n$ and using equations (8.2), we get $\alpha_1 = 0, \alpha_2 = 0, \dots, \alpha_n = 0$. Consequently, the vectors $e^1, e^2, \dots, e^n$ are linearly independent, i.e. form a basis.

Thus the reciprocal basis e^j of the basis e_i exists and is defined uniquely.

Remark 1. By virtue of the symmetry of relations (8.2) with respect to e^j and e_i the basis e_i is a reciprocal basis of e^j . Therefore, in what follows we shall speak of the reciprocal bases e_i and e^j .

Remark 2. If the basis $e_1, e_2, \dots, e_n$ is orthonormal, then the reciprocal basis e^j coincides with the given basis. Indeed, setting $e^j = e_j$ in this case, we make sure that relations (8.2) are satisfied. Using the property of uniqueness of the reciprocal basis, we verify the validity of the remark.

Assume that e_i, e^j are reciprocal bases and x is an arbitrary vector of the space.

Resolving the vector x into the components e_i and e^j , we get

$$\left. \begin{aligned} x &= x_1 e^1 + x_2 e^2 + \dots + x_n e^n, \\ x &= x^1 e_1 + x^2 e_2 + \dots + x^n e_n. \end{aligned} \right\} \quad (8.5)$$

The coordinates $(x_1, x_2, \dots, x_n)$ of the vector x in the basis e^j are said to be **covariant coordinates** of the vector x , and the coordinates $(x^1, x^2, \dots, x^n)$ of that vector in the basis e_i are said to be **contravariant coordinates** of the vector x . We shall explain these terminology in 8.1.3.

For example,

The agreement on summation makes it possible to write formulas (8.5) in the following compact form: 8

Remark 3. Similar superscripts and subscripts which were mentioned in the agreement on summation are usually called the sum indices. It is clear that we can denote the sum indices by any identical symbols without changing the expression in which they appear. For instance, $x_i e^i$ and $x_j e^j$ constitute the same expression.

279

For that purpose, we multiply scalarly the first equation. (8.6) by $\mathbf{e}_j$ and the second equation by $\mathbf{e}^j$. Taking into account relations (8.2), we find that

$$(\mathbf{x}, \mathbf{e}_i) = x_i(\mathbf{e}^i, \mathbf{e}_i) = \delta_i^i = x_j, (\mathbf{x}, \mathbf{e}^j) = x^j(\mathbf{e}_i, \mathbf{e}^j) = x^i\delta_i^j = x^j.$$

Thus we have

$$x_j(\mathbf{x}, \mathbf{e}_j), x^j = (\mathbf{x}, \mathbf{e}^j). \quad (8.7)$$

Using relations (8.7), we write formulas (8.6) as

$$\mathbf{x} = (\mathbf{x}, \mathbf{e}_i) \mathbf{e}^i, \mathbf{x} = (\mathbf{x}, \mathbf{e}^i) \mathbf{e}_i. \quad (8.8)$$

Relations (8.8) are known as the **Gibbs formulas**.

Let us again direct our attention to the problem of constructing reciprocal bases.

With the aid of formulas (8.8) we get

$$\mathbf{e}^i = (\mathbf{e}^i, \mathbf{e}^j) \mathbf{e}_j, \mathbf{e}_i = (\mathbf{e}_i, \mathbf{e}^j) \mathbf{e}^j. \quad (8.9)$$

We introduce the designations

$$g_{ij} = (\mathbf{e}_i, \mathbf{e}_j), g^{ij} = (\mathbf{e}^i, \mathbf{e}^j). \quad (8.10)$$

Using these designations, we rewrite relations (8.9) as

$$\mathbf{e}^i = g^{ij} \mathbf{e}_j, \mathbf{e}_i = g_{ij} \mathbf{e}^j. \quad (8.11)$$

Thus, to construct the basis $\mathbf{e}^i$ from the basis $\mathbf{e}_i$, it is sufficient to know the matrix (g^{ij}) , and to construct the basis $\mathbf{e}_i$ from the basis $\mathbf{e}^i$, it is sufficient to know the matrix (g_{ij}) .

Let us prove that the indicated matrices are mutually inverse. Note that since the elements of the inverse matrix can be calculated in terms of the elements of the given matrix, it is clear that the problem of constructing reciprocal bases can be solved with the aid of relations (8.11).

Let us ascertain that the matrices (g^{ij}) and (g_{ij}) are mutually inverse.

Multiplying scalarly the first equation (8.11) by e_k , we get $(e^i, e_k) g^{ij} (e_j, e_k)$.

Taking (8.2) and (8.10) into account, we find from this relation that

$$g^{ij} g_{jk} = \delta_k^i = \begin{cases} 1 & \text{for } i = j, \\ 0 & \text{for } i \neq j. \end{cases}$$

Thus, the product of the matrices (g^{ij}) and (g_{ij}) is an identity matrix.

Consequently, the matrices (g^{ij}) and (g_{ij}) are mutually inverse.

8.1.3. Transformation of a basis and coordinates. Suppose e_i and e^i are given reciprocal bases and e_i and $e^{i'}$ are some new reciprocal bases whose elements are denoted by primed indices. In fact this means that we introduce a new natural series $1', 2', 3', \dots$ and suppose that the index i' assumes the values $1', 2', \dots, n'$. Thus the indices i and i' independently assume different values

$$i = 1, 2, \dots, n, \quad 1', 2', \dots, n'.$$

Using the agreement on summation introduced in 8.1.2, we write the formulas for the transformation of base vectors. As a result we get :

(1) formulas for transition from the old basis e_i to a new basis $e^{i'}$ and formulas for a reverse transition

$$e_{i'} = b_i^{i'} e_i, \quad e^{i'} = \tilde{b}_i^{i'} e^i, \quad i = 1, 2, \dots, n, \quad i' = 1', 2', \dots, n', \quad (8.12)$$

(2) formulas for transition from the old basis e^i to a new basis $e^{i'}$ and formulas for a reverse transition:

$$e^{i'} = \tilde{b}_i^{i'} e^i, \quad e^i = \tilde{b}_i^{i'} e^{i'}, \quad i = 1, 2, \dots, n, \quad i' = 1', 2', \dots, n'. \quad (8.13)$$

Since transformations (8.12) (as well as (8.13)) are mutually reciprocal, the matrices $(b_i^{i'})$ and $(\tilde{b}_i^{i'})$ (as well as $(\tilde{b}_i^{i'})$ and $(b_i^{i'})$) are mutually inverse.

Let us prove that the matrices $(b_i^{i'})$ and $(\tilde{b}_i^{i'})$ are identical. We shall thus prove that the matrices $(b_i^{i'})$ and $(\tilde{b}_i^{i'})$ are identical. To prove this, we multiply scalarly the first equation (8.12) by e^k and the second equation (8.13) by $e_{k'}$. Taking relations (8.2) into account, we obtain

$$\begin{aligned} (e_{i'}, e^k) &= b_{i'}^i (e_i, e^k) = b_{i'}^i \delta_i^k = b_{i'}^i, \\ (e^i, e_{k'}) &= \tilde{b}_{i'}^i (e^{i'}, e_{k'}) = \tilde{b}_{i'}^i \delta_{i'}^{k'} = \tilde{b}_{i'}^i. \end{aligned}$$

From these relations we get for $k = i, k' = i'$ the following equations :

$$b_{i'}^i = (e_{i'}, e^i), \quad (8.14)$$

$$\tilde{b}_{i'}^i = (e^i, e_{i'}). \quad (8.15)$$

Since the right-hand sides of relations (8.14) and (8.15) are equal, the left-hand sides are equal too. In other words, $b_{i'}^i = \tilde{b}_{i'}^i$ and this means that the matrices $(b_{i'}^i)$ and $(\tilde{b}_{i'}^i)$ are identical. Note that the elements $b_{i'}^i$ of the matrix $(b_{i'}^i)$ can be calculated by formulas (8.14).

Thus, the following statement holds true.

To pass from the basis e_i, e^i to the basis $e_{i'}, e^{i'}$, it is sufficient to know the matrix $(b_{i'}^i)$ of transition from the basis e_i to the basis $e_{i'}$ (the matrix $(b_{i'}^i)$ can be calculated from the (matrix $(b_{i'}^i)$).

Here is a full collection of formulas for the transformation of base vectors:

$$\left. \begin{aligned} e_{i'} &= b_{i'}^i e_i, & e_i &= b_i^{i'} e_{i'}, \\ e^{i'} &= b_i^{i'} e^i, & e^i &= b_{i'}^i e^{i'}. \end{aligned} \right\} \quad (8.16)$$

Let us derive formulas for the transformation of the coordinates of the vector x in passing to a new basis.

Let $x_{i'}$ be the covariant coordinates of the vector $\mathbf{x}$ in the basis $\mathbf{e}_{i'}$, $\mathbf{e}^{i'}$. Then, according to (8.7), we have $x_{i'} = (\mathbf{x}, \mathbf{e}_{i'})$. Substituting the expression for $\mathbf{e}_{i'}$ from formulas (8.16) into the right-hand side of this relation, we find that $x_{i'} = (\mathbf{x}, b_{i'}^j \mathbf{e}_j) = b_{i'}^j (\mathbf{x}, \mathbf{e}_j) = b_{i'}^j x_j$.

We arrive at the following conclusion : *when we pass to a new basis the formulas for the transformation of the covariant coordinates the vector $\mathbf{x}$ have the form*

$$x_{i'} = b_{i'}^j x_j. \quad (8.17)$$

Consequently, *when we pass to a new basis, the covariant coordinates of the vector $\mathbf{x}$ are transformed by means of the matrix $(b_{i'}^j)$ of a direct transition from the old basis to a new one.*

This coherence of transformations explains the name-of the *covariant** coordinates of a vector.

Let us now consider the transformation of the contravariant coordinates of the vector $\mathbf{x}$.

Substituting the expression for $\mathbf{e}^{i'}$ from formulas (8.16) into the right-hand side of the relation $x^{i'} = (\mathbf{x}, \mathbf{e}^{i'})$, we find, after some transformations, that

$$x^{i'} = b_i^{i'} x_i. \quad (8.18)$$

We see that *when we pass to a new basis, the contravariant coordinates of the vector $\mathbf{x}$ are transformed by means of the matrix $(b_i^{i'})$ of the reverse transition from the new basis to the old one. This non concordance of the transformations explains the name of the contravariant coordinates of a vector.*

8.2. The Concept of a Tensor. Principal Operations on Tensors

8.2.1. The concept of a tensor. We discuss here an arbitrary (not necessarily Euclidean) real n -dimensional linear space L^n .

Definition. A *tensor A of the type (p, q)* (p times covariant and q times contravariant) is a geometrical object which (1) is defined by the $n^p + n^q$

coordinates $A_{i_1 \dots i_p}^{k_1 \dots k_q}$ in every basis $\mathbf{e}_i$ of the linear space L^n (the indices $i_1, \dots, i_p, k_1, \dots, k_q$ independently assume the values $1, 2, \dots, n$), (2), possesses the property that its coordinates $A_{i'_1 \dots i'_p}^{k'_1 \dots k'_q}$ the basis $\mathbf{e}_{i'}$ are connected with the coordinates $A_{i_1 \dots i_p}^{k_1 \dots k_q}$ in the basis $\mathbf{e}_i$ by the relations

$$A_{i'_1 \dots i'_p}^{k'_1 \dots k'_q} = b_{i'_1}^{i_1} \dots b_{i'_p}^{i_p} b_{k_1}^{k'_1} \dots b_{k_q}^{k'_q} A_{i_1 \dots i_p}^{k_1 \dots k_q}. \quad (8.19)$$

in which $b_{i'}^i$ are the elements of the matrix $(b_{i'}^i)$ of the transition from the basis $\mathbf{e}_i$ to the basis $\mathbf{e}_{i'}$ and $b_k^{k'}$ are the elements of the matrix of the reverse transition from $\mathbf{e}_{i'}$ to $\mathbf{e}_i$.

* Covariant means concordantly changing.

The number $r = p + q$ is called the rank of the tensor.

Remark 1. Formulas (8.19) are called the formulas for the transformation of the coordinates of a tensor in the transformation of a basis.

Note that the covariant and contravariant coordinates of a vector are transformed by formulas (8.19) (for $p = 1$ and $q = 0$ in the first case and for $p = 0$ and $q = 1$ in the second case, see 8.1.3.). Therefore, a vector is a tensor of rank 1 (once covariant or once contravariant depending on the choice of the type of coordinates of the vector).

Note that there are also *tensors of rank 0*. These are tensors having only one coordinate, the coordinate having no indices and being of the same value in all the coordinate systems. The tensors of rank 0 are usually called **invariant tensors**.

Remark 2. The indices $i_1, \dots, i_p$ are said to be covariant and $k_1, \dots, k_q$, contravariant. The name can be explained by the fact that with respect of each index the transformation of the coordinates of a tensor is carried out by a complete analogy with the transformations of the covariant and contravariant coordinates of a vector (see formulas (8.17) and (8.18)).

To make the definition of a tensor correct, we must make sure that the successive transitions from the basis e_i to the basis $e_{i'}$ and then from $e_{i'}$ to $e_{i''}$ can be reduced to the same transformation of the coordinates of a tensor as in the direct transition from e_i to $e_{i''}$.

Let $(b_{i'}^i)$, $(b_{i''}^{i'})$ and $(b_{i''}^i)$ be, respectively, the matrices of the transition from the basis e_i to the basis $e_{i'}$, from $e_{i'}$ to $e_{i''}$ and from e_i to $e_{i''}$. Since the transformation matrices are multiplied in successive transitions, the following relations are obvious:

$$b_{i''}^i = b_{i''}^{i'} b_{i'}^i, \quad b_{i'}^{i''} = b_{i'}^i b_{i''}^{i'}. \quad (8.20)$$

After the remarks we have made, let us verify the correctness of the definition of a tensor. Let $A_{i_1 \dots i_p}^{k_1 \dots k_q}$, $A_{i'_1 \dots i'_p}^{k'_1 \dots k'_q}$, $A_{i''_1 \dots i''_p}^{k''_1 \dots k''_q}$ be the coordinates of the tensor A in the basis e_i , $e_{i'}$ and $e_{i''}$ respectively. By formulas (8.19) for the successive transition from e_i to $e_{i'}$ and then to $e_{i''}$ we get

$$A_{i'_1 \dots i'_p}^{k'_1 \dots k'_q} = b_{i'_1}^{i_1} \dots b_{i'_p}^{i_p} b_{k'_q}^{k_q} \dots b_{k'_1}^{k_1} A_{i_1 \dots i_p}^{k_1 \dots k_q}. \quad (8.21)$$

$$A_{i''_1 \dots i''_p}^{k''_1 \dots k''_q} = b_{i''_1}^{i'_1} \dots b_{i''_p}^{i'_p} b_{k''_q}^{k'_q} \dots b_{k''_1}^{k'_1} A_{i'_1 \dots i'_p}^{k'_1 \dots k'_q}. \quad (8.22)$$

Substituting the expression $A_{i'_1 \dots i'_p}^{k'_1 \dots k'_q}$ for the coordinates from (8.21) and taking relations (8.20) into account, we obtain

$$\begin{aligned} A_{i''_1 \dots i''_p}^{k''_1 \dots k''_q} &= \left(b_{i''_1}^{i'_1} b_{i''_1}^{i_1} \right) \dots \left(b_{i''_p}^{i'_p} b_{i''_p}^{i_p} \right) \left(b_{k''_1}^{k'_1} b_{k''_1}^{k_1} \right) \\ &\dots \left(b_{k''_q}^{k'_q} b_{k''_q}^{k_q} \right) A_{i'_1 \dots i'_p}^{k'_1 \dots k'_q} = b_{i''_1}^{i_1} \dots b_{i''_p}^{i_p} b_{k''_q}^{k_q} \dots b_{k''_1}^{k_1} A_{i_1 \dots i_p}^{k_1 \dots k_q}. \end{aligned}$$

Thus consecutive transitions from the basis e_i to the basis $e_{i'}$ and then to the basis $e_{i''}$ lead to the same transformation of the coordinates of a tensor as in a direct

transition from e_i to $e_{i''}$. We have established the correctness of the definition of a tensor.

Remark 3. Any system n^{p+n} of the numbers $A_{i_1 \dots i_p}^{k_1 \dots k_q}$ can be regarded in the given basis e_i as the coordinates of a certain tensor A of type (p, q) . To verify this fact, we define in an arbitrary basis $e_{i'}$ by means of formulas (8.19), a system of numbers $A_{i'_1 \dots i'_p}^{k'_1 \dots k'_q}$ which will be regarded as the coordinates of the required tensor A in the basis $e_{i'}$. These coordinates are evidently transformed by formulas (8.19) when we pass from the basis e_i to the basis $e_{i'}$. As before, it is easy to verify that the consecutive transitions from the basis e_i to the basis $e_{i'}$ and then to the basis $e_{i''}$ lead to the same transformation of the coordinates obtained as in a direct transition from e_i to $e_{i''}$. Consequently, the system of numbers $A_{i_1 \dots i_p}^{k_1 \dots k_q}$ indeed constitutes the coordinates of a certain tensor A of type (p, q) .

8.2.2. Examples of tensors.

1°. **A null-tensor.** Among the tensors of type (p, q) we must isolate the so-called *null-tensor*. This is a tensor whose coordinates in any basis are zero. Relations (8.19) are evidently satisfied.

Note that if the coordinates of the tensor A are zero in some basis, then, according to (8.19), they are zero in any basis and, consequently, A is a null-tensor.

2°. **The Kronecker delta.** Let us verify that the tensor A of type $(1, 1)$, having the coordinates δ_i^k in the basis e_i , has the coordinates $\delta_{i'}^{k'}$ in the basis $e_{i'}$.

Thus, let A be a tensor having the coordinates δ_i^k in the basis e_i . To find the coordinates of this tensor in the basis $e_{i'}$, we must make use of formulas (8.19), i.e. the coordinates of the tensor A in the basis $e_{i'}$, are equal to $b_k^{k'} b_i^i \delta_i^k$. Using the properties of the Kronecker delta, we find that $b_k^{k'} b_i^i \delta_i^k = b_k^{k'} b_i^k = \delta_{i'}^{k'}$.

Thus, in the new basis $e_{i'}$ the coordinates of the tensor A are indeed equal to $\delta_{i'}^{k'}$.

We can, therefore, regard the Kronecker delta as a tensor of type (1, 1).

3°. Suppose $A(x, y)$ is a bilinear form defined in a finite-dimensional Euclidean space E^n , and $e_1, e_2, \dots, e_n$ is some basis in that space, Then the vectors x and y can be represented in the form $x = x^i e_i, y = y^j e_j$.

Using the linear property of the form $A(x, y)$ with respect to each argument, we can write

$$A(x, y) = A(x^i e_i, y^j e_j) = A(e_i, e_j) x^i y^j.$$

Designating $A(e_i, e_j)$ as a_{ij} , we get

$$a_{ij} = A(e_i, e_j). \quad (8.23)$$

Then the form $A(x, y)$ can be written as follows:

$$A(e_i, e_j) = a_{ij} x^i y^j. \quad (8.24)$$

Let us ascertain that the coefficients a_{ij} of the matrix having the form $A(x, y)$ are transformed by the law (8.19) of transformation of the coordinates of a tensor of type (2, 0) when we pass to a new basis, i.e. they constitute a tensor of type (2, 0).

Let us consider an arbitrary basis $e_{1'}, e_{2'}, \dots, e_{n'}$. We write the form $A(x, y)$ in this basis in the form (8.24), i.e. $A(x, y) = a_{i'j'} x^{i'} y^{j'}$, with

$$a_{i'j'} = A(e_{i'}, e_{j'}). \quad (8.25)$$

We pass from the basis $e_1, e_2, \dots, e_n$ to a new basis $e_{1'}, e_{2'}, \dots, e_{n'}$. Designating the transition matrix from the basis e_i to the basis $e_{i'}$ as $b_{i'}^i$, we get

$$e_{i'} = b_{i'}^i e_i, e_{j'} = b_{j'}^j e_j.$$

Substituting these expressions for $\mathbf{e}_{i'}$ and $\mathbf{e}_{j'}$ into the right-hand side of (8.25) and using the linear property of the form $A(\mathbf{x}, \mathbf{y})$ with respect to each argument, we find that $a_{i'j'} = A(b_{i'}^i \mathbf{e}_i, b_{j'}^j \mathbf{e}_j) = b_{i'}^i b_{j'}^j A(\mathbf{e}_i, \mathbf{e}_j)$.

In accordance with formula (8.23), the last relation can be rewritten as $a_{i'j'} = b_{i'}^i b_{j'}^j a_{ij}$.

Consequently, the coefficients of the matrix of a bilinear form are transformed by the law (8.19) of transformation of the coordinates of a tensor of type (2,0) and can, therefore, be regarded as the coordinates of a tensor of that type.

4°. Every linear operator given in a finite-dimensional Euclidean space E^n and acting into the same space can be put into correspondence with a certain tensor of type (1, 1), and this tensor completely defines the indicated operator.

Suppose $\mathbf{y} = L\mathbf{x}$ is a linear operator defined in E^n , and $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ is a basis in E^n . Since $\mathbf{x} = x^i \mathbf{e}_i$ and $\mathbf{y} = y^j \mathbf{e}_j$ and L is a linear operator, we have

$$y^j \mathbf{e}_j = x^i L(\mathbf{e}_i). \quad (8.26)$$

We resolve the vector $L(\mathbf{e}_i)$ into the components $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$.

$$L(\mathbf{e}_i) = a_i^j \mathbf{e}_j.$$

Substituting the expression obtained for $L(\mathbf{e}_i)$ into (8.26) and using the uniqueness of resolution into components, we get

$$y^j = a_i^j x^i, \quad j = 1, 2, \dots, n. \quad (8.27)$$

Recall that relations (8.27) can be regarded as the coordinate method of defining a linear operator. Then the matrix (a_i^j) of the coefficients a_i^j is called the matrix of a linear operator.

Let us ascertain that the coefficients of this matrix are transformed by the law (8.19) of transformation of the coordinates of a tensor of type (1, 1) when we pass to a new basis and, therefore, constitute a tensor of type (1, 1).

Let us consider an arbitrary basis $\mathbf{e}_{1'}, \mathbf{e}_{2'}, \dots, \mathbf{e}_{n'}$. We write the linear operator L in this basis in the form (8.27):

$$y^{j'} = a_{i'}^{j'} x^{i'}, \quad j' = 1', 2', \dots, n'. \quad (8.28)$$

We pass from the basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ to the basis $\mathbf{e}_{1'}, \mathbf{e}_{2'}, \dots, \mathbf{e}_{n'}$. Designating the transition matrix $(b_{i'}^j)$ (or, what is the same, $(b_{j'}^j)$), we get* (see 8.1.3)

(* We designate the sum index in the formula for y^j as k' .)

$$x^i = b_{i'}^i x^{i'}, \quad y^j = b_{k'}^j y^{k'}.$$

We substitute these expressions for x^i and y^j into (8.27) and get the following relations: _

$$y^{k'} b_{k'}^i = a_{i'}^j b_{i'}^i x^{i'}, \quad j = 1, 2, \dots, n. \quad (8.29)$$

We must get expression for $y^{j'}$ from (8.29). For that purpose, we multiply both sides of (8.29) by $b_j^{j'}$ and perform a summation by j from 1 to n . Taking into consideration that $b_{k'}^i b_j^{j'} = \delta_{k'}^{j'}$, we get $y^{k'} \delta_{k'}^{j'} = (b_{i'}^i b_j^{j'} a_{i'}^j) x^{i'}$. Note that $y^{k'} \delta_{k'}^{j'} = y^{j'}$. Therefore $y^{j'} = (b_{i'}^i b_j^{j'} a_{i'}^j) x^{i'}$. Comparing this expression for $y^{j'}$ with the expression for $y^{j'}$ obtained by formula (8.28), we get the following identity which is valid for any vectors $\mathbf{x}$ (for any coordinates $x^{i'}$):

$$a_{i'}^{j'} x^{i'} = (b_{i'}^i b_j^{j'} a_{i'}^j) x^{i'}.$$

It follows from this and from the arbitrariness of $x^{i'}$ that the coefficients at of the matrix of a linear operator are transformed by the law $a_{i'}^{j'} = b_{i'}^i b_j^{j'} a_i^j$.

Thus, the coefficients a_i^j are transformed by the law (8.19) of transformation of the coordinates of a tensor of type (1, 1) and, therefore, constitute such a tensor.

8.2.3. The principal operations on tensors. The principal operations on tensors are the *addition and subtraction of tensors, multiplication of a tensor by a number, multiplication of tensors, contraction of tensors, commutation of indices, symmetrisation and alternation of tensors.*

Let us define these operations.

1°. Addition and subtraction of tensors. The operations of addition and subtraction are defined for tensors of the same type.

Suppose A and B are two tensors of the same type (p, q) , and $A_{i_1 \dots i_p}^{k_1 \dots k_q}$ and $B_{i_1 \dots i_p}^{k_1 \dots k_q}$ are the coordinates with like indices of these tensors in the basis e_i .

The **sum** $(A + B)$ (the **difference** $A - B$) of these tensors is a tensor having the coordinates

$$A_{i_1 \dots i_p}^{k_1 \dots k_q} + B_{i_1 \dots i_p}^{k_1 \dots k_q} \left(A_{i_1 \dots i_p}^{k_1 \dots k_q} - B_{i_1 \dots i_p}^{k_1 \dots k_q} \right)$$

in the basis e_i .

For the definition of the operations of addition and subtraction of tensors to be correct, we must verify whether the coordinates of the sum (the difference) of the tensors are transformed by the law (8.19) of transformation of the coordinates of a tensor. To do that, we assume that alongside formulas (8.19) for the transformation of the coordinates of the tensor A we have analogous formulas for the transformation of the coordinates of the tensor B . Then by adding (subtracting) two formulas of this kind for the transformation of the coordinates of the tensors A and B we ascertain that the coordinates of the sum (the difference) are transformed by the law of transformation of the coordinates of a tensor. .

2°. Multiplication of a tensor by a number. Suppose A is a tensor of type (p, q) having the coordinates $A_{i_1 \dots i_p}^{k_1 \dots k_q}$ in the basis e_i and α is an arbitrary real number.

The **product** αA of the tensor A by the number α is a tensor having the coordinates $\alpha A_{i_1 \dots i_p}^{k_1 \dots k_q}$ in the basis e_i .

The fact that the coordinates $\alpha A_{i_1 \dots i_p}^{k_1 \dots k_q}$ are transformed by the law of tensors, can be seen from formulas (8.19).

3°. Multiplication of tensors. The operation of multiplication of tensors is defined for tensors of arbitrary type.

Suppose A is a tensor of type (p, q) having coordinates $A_{i_1 \dots i_p}^{k_1 \dots k_q}$ in the given basis e_i , and B is a tensor of type (r, s) having the coordinates $B_{l_1 \dots l_r}^{m_1 \dots m_s}$ in the same basis.

To define the **product** $D = AB$ of the tensors A and B , we form various products of the coordinates of the tensor A by the coordinates of the tensor B . In every such product, the indices $l_1, \dots, l_r$ and $m_1, \dots, m_s$ of the coordinates of the tensor B are replaced by new indices. Namely, we set $l_1 = i_{p+1}, \dots, l_r = i_{p+r}$ and $m_1 = k_{q+1}, \dots, m_s = k_{q+s}$.

The product $D = AB$ of the tensors A and B is a tensor of type $(p+r, q+s)$, having the coordinates

$$D_{i_1 \dots i_p i_{p+1} \dots i_{p+r}}^{k_1 \dots k_q k_{q+1} \dots k_{q+s}} = A_{i_1 \dots i_p}^{k_1 \dots k_q} B_{i_{p+1} \dots i_{p+r}}^{k_{q+1} \dots k_{q+s}} \quad (8.30)$$

in the basis e_i .

To verify that the coordinates $D_{i_1 \dots i_p i_{p+1} \dots i_{p+r}}^{k_1 \dots k_q k_{q+1} \dots k_{q+s}}$ defined by relation (8.30) are transformed by the law of transformation of the coordinates of a tensor when we pass to a new basis, i.e. that they are indeed the coordinates of a tensor, it is sufficient to write the transformation formulas (8.19) for the coordinates $A_{i_1 \dots i_p}^{k_1 \dots k_q}$ and

$B_{i_{p+1} \dots i_{p+r}}^{k_{q+1} \dots k_{q+s}}$ of the tensors A and B and multiply the right-hand and left-hand sides of the formulas. As a result, using relations (8.30), we can easily obtain the necessary formulas for the transformation of the coordinates $D_{i_i \dots i_p i_{p+1} \dots i_{p+r}}^{k_1 \dots k_q k_{q+1} \dots k_{q+s}}$.

Remark. The operation of multiplication of tensors does not possess the commutative property: in general $AB \neq BA$. This can be explained by the fact that the sequence of the indices in the coordinates of a tensor defines the “number” of that coordinate.

Thus, although the numerical value of the expressions

$$A_{i_1 \dots i_p}^{k_1 \dots k_q} B_{i_{p+1} \dots i_{p+r}}^{k_{q+1} \dots k_{q+s}} \text{ and } B_{i_{p+1} \dots i_{p+r}}^{k_{q+1} \dots k_{q+s}} A_{i_1 \dots i_p}^{k_1 \dots k_q}$$

is the same, the sequence of indices in these expressions is different and, therefore, they correspond to the coordinates with different “numbers”. And this means that $AB \neq BA$.

4°. **Contraction of a tensor.** The operation of **contraction** is applied to a tensor of type (p, q) in which $p \neq 0$ and $q \neq 0$ (i.e. to a tensor which has at least one super index and one subscript).

Let A be a tensor of the indicated type. We shall describe the operation of contraction.

The contraction of the tensor A is performed by some marked superscript and subscript. As a result of contraction, a tensor of type $(p-1, q-1)$ is obtained.

Assume, for instance, that a superscript with the number m and a subscript with the number n are marked in each of the coordinates of the tensor A :

$$A_{i_1 \dots i_n \dots i_p}^{k_1 \dots k_m \dots k_q}.$$

Let us sum (contract) the coordinates of the tensor with similar marked indices. This operation yields the following coordinates of the tensor obtained by a contraction with respect to the superscript and the subscript with the numbers m and n respectively:

$$A_{i_1 \dots \alpha \dots i_p}^{k_1 \dots \alpha \dots k_q} \quad (8.31)$$

(in expression (8.31) we have used the agreement on summation).

Let us verify that quantities (8.31) indeed form the coordinates of a tensor of type $(p-1, q-1)$. For that purpose we direct our attention to formulas (8.19). We rewrite these formulas in the following form, having marked the indices of interest to us (we under-line these indices by brackets) :

$$A_{i_1' \dots \underline{i_m'} \dots i_p'}^{k_1' \dots \underline{k_m'} \dots k_q'} = b_{k_1}^{k_1'} \dots b_{k_m}^{\underline{k_m'}} \dots b_{k_q}^{k_q'} b_{i_1'}^{i_1'} \dots b_{i_n'}^{\underline{i_n'}} \dots b_{i_p'}^{i_p'} A_{i_1' \dots i_n' \dots i_p'}^{k_1' \dots k_m' \dots k_q'}. \quad (8.32)$$

Let us perform a summation in the right-hand and left-hand sides of (8.32) with respect to the emphasized indices k_m' and i_n' . To do that, it is sufficient to set these indices equal to α' and make use of the agreement on summation. As a result, we get a certain equality. We find the expressions on the right-hand and left-hand sides of this equality. On the left-hand side we evidently get an expression

$$A_{i_1' \dots \alpha' \dots i_p'}^{k_1' \dots \alpha' \dots k_q'}. \quad (8.33)$$

On the right-hand side the product of $b_{k_m}^{\alpha'}$ by $b_{\alpha'}^{i_n'}$ is equal to $\delta_{k_m}^{i_n'}$, i.e. equal to unity for $k_m = i_n = \alpha$ and to zero for $k_m \neq i_n$. Thus, on the right-hand side we get an expression

$$b_{k_1}^{k_1'} \dots b_{k_q}^{k_q'} b_{i_1'}^{i_1'} \dots b_{i_p'}^{i_p'} A_{i_1' \dots \alpha \dots i_p'}^{k_1' \dots \alpha \dots k_q'}. \quad (8.34)$$

Comparing expressions (8.33) and (8.34), we make sure that the quantities $A_{i_1 \dots i_p}^{k_1 \dots k_q}$ are transformed in passing to a new basis by the law of transformation of the coordinates of a tensor. This tensor is evidently of the type $(p - 1, q - 1)$.

Remark. The term “contraction of tensors” can also be used in the following sense.

Let us consider two tensors A and B . The coordinates of the first tensor have at least one superscript k and the coordinates of the second have at least one subscript i .

We form a product AB of these tensors and then contract the tensor AB with respect to the superscript k and the subscript i . The terminology usually used for this operation is the “*contraction of tensors A and B with respect to the indices k and i* ”.

5°. **Commutation of indices.** This operation consists in subjecting the indices of each coordinate of a tensor to the same commutation in any basis. This means that we “enumerate” the coordinates of a given tensor in a different way. The reader is invited to verify that a tensor (in general different from the given tensor) results.

6°. **Symmetrisation and alternation.** We must first introduce the concepts of a **symmetric** and a **skew-symmetric** tensor.

The tensor A with the coordinates

$$A_{i_1 \dots i_m \dots i_n \dots i_p}^{k_1 \dots k_q} \quad (8.35)$$

is said to be **symmetric** with respect to the subscripts i_m and i_n if as a result of the commutation of these indices* the coordinates of the tensor A do not change their meaning, i.e.

$$A_{i_1 \dots i_m \dots i_n \dots i_p}^{k_1 \dots k_q} = A_{i_1 \dots i_n \dots i_m \dots i_p}^{k_1 \dots k_q}. \quad (8.36)$$

Relation (8.36) is called the **condition of symmetry** of the tensor A with respect to subscripts labelled m and n .

The tensor A is said to be **skew-symmetric** with respect to the subscripts i_m and i_n if upon an interchange of these indices the following relation holds true:

$$A_{i_1 \dots i_m \dots i_n \dots i_p}^{k_1 \dots k_q} = -A_{i_1 \dots i_n \dots i_m \dots i_p}^{k_1 \dots k_q}. \quad (8.37)$$

Relation (8.37) is called the **condition of skew-symmetry** of the tensor A with respect to the subscripts labelled m and n .

The concepts of symmetry and skew-symmetry of a tensor with respect to two superscripts can be introduced by analogy.

Remark. If the condition of symmetry (skew-symmetry) with respect to two subscripts i_m and i_n is fulfilled for the tensor A in a given system of coordinates, then it is also fulfilled in any other system of coordinates.

We shall now describe the **operation of symmetrization**.

Let A be a tensor of type (p, q) with coordinates (8.35). We inter-change the subscripts with the numbers m and n in each coordinate and construct a tensor $A_{(m, n)}$ with coordinates

$$\frac{1}{2} \left(A_{i_1 \dots i_m \dots i_n \dots i_p}^{k_1 \dots k_q} + A_{i_1 \dots i_n \dots i_m \dots i_p}^{k_1 \dots k_q} \right). \quad (8.38)$$

The operation of constructing a tensor $A_{(m, n)}$ is called the **operation of symmetrisation** of the tensor A with respect to the subscripts labelled m and n .

Note that coordinates (8.38) of the tensor $A_{(m, n)}$ are usually symbolized as

$$A_{i_1 \dots (i_m \dots i_n) \dots i_p}^{k_1 \dots k_q}. \quad (8.39)$$

The condition of symmetry (8.36) with respect to the subscripts labelled m and n is evidently fulfilled for the tensor $A_{(m, n)}$.

The operation of symmetrisation of a tensor **with respect to the superscripts** labelled m and n can be defined by analogy. The tensor that is constructed is symbolized as $A^{(m, n)}$. The designation used for the coordinates of the tensor $A^{(m, n)}$ is similar to (8.39) for the coordinates of the tensor $A_{(m, n)}$.

The **operation of alternating** the tensor A with respect to the subscripts labelled m and n is carried out as follows.

* Recall that when we commute the indices of the coordinates of a tensor, we obtain in general coordinates of another tensor.

The subscripts labelled m and n of each coordinate of the tensor A are interchanged and then a tensor $A_{(m, n)}$ with coordinates

$$\frac{1}{2} \left(A_{i_1 \dots i_m \dots i_n \dots i_p}^{k_1 \dots k_q} - A_{i_1 \dots i_n \dots i_m \dots i_p}^{k_1 \dots k_q} \right). \quad (8.40)$$

is constructed.

The operation of constructing a tensor $A_{(m, n)}$ is called the **operation of alternation** of the tensor A with respect to the subscripts labelled m and n . Coordinates (8.40) of the tensor $A_{[m, n]}$ are usually symbolized as

$$A_{i_1 \dots [i_m \dots i_n] \dots i_p}^{k_1 \dots k_q}. \quad (8.41)$$

Condition (8.37) of skew-symmetry with respect to the subscripts i_m and i_n is evidently fulfilled for the tensor $A_{[m, n]}$.

The operation of alternating a tensor **with respect to the superscripts** i_m and i_n can be carried out by analogy. The constructed tensor is symbolized as $A^{[m, n]}$. The designation similar to (8.41) for the coordinates of the tensor $A_{[m, n]}$ is used for the coordinates of the tensor $A^{[m, n]}$.

We must point out in conclusion the obvious equality

$$A = A_{(m, n)} + A_{[m, n]}.$$

8.3. A Metric Tensor. Principal Operations of Vector Algebra in Tensor Designations

8.3.1. The concept of a metric tensor in a Euclidean space. We mentioned in 7.2 that in a finite-dimensional linear space a scalar product could be defined by means

of a bilinear form which is a polar of some positive-definite quadratic form. We assume in this section that in the finite-dimensional Euclidean space E^n the scalar product is defined by this type of a bilinear form. $A(\mathbf{x}, \mathbf{y})$.

We have ascertained in 8.2.2 (Example 3) that the coefficients of the matrix of a bilinear form can be regarded as the coordinates of a tensor. For the bilinear form $A(\mathbf{x}, \mathbf{y})$ which defines the scalar product in E^n , we shall designate these coefficients as g_{ij} . Thus g_{ij} are the coordinates of the tensor G in the basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$. This tensor of type $(2, 0)$ is called a **metric tensor** of the space E^n . Recall that the coordinates g_{ij} of the tensor G are defined by the relations

$$g_{ij} = A(\mathbf{e}_i, \mathbf{e}_j) \quad (8.42)$$

(see formula (8.28)).

It should also be pointed out that since the form $A(\mathbf{x}, \mathbf{y})$ is symmetric ($A(\mathbf{x}, \mathbf{y}) = A(\mathbf{y}, \mathbf{x})$), it follows, according to (8.42), that $g_{ij} = g_{ji}$, i.e. the *metric tensor* G is *symmetric with respect to the subscripts* i and j .

Suppose that $\mathbf{x}$ and $\mathbf{y}$ are arbitrary vectors in E^n , and x^i and y^j are the coordinates of these vectors in the basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$.

The scalar product $(\mathbf{x}, \mathbf{y})$ of the vectors $\mathbf{x}$ and $\mathbf{y}$ is equal to $A(\mathbf{x}, \mathbf{y})$. Directing our attention to expression (8.24) for a bilinear form in the given basis and using the equation $(\mathbf{x}, \mathbf{y}) = A(\mathbf{x}, \mathbf{y})$, we get the following formula for the scalar product $(\mathbf{x}, \mathbf{y})$ of the vectors $\mathbf{x}$ and $\mathbf{y}$:

$$(\mathbf{x}, \mathbf{y}) = g_{ij} x^i y^j. \quad (8.43)$$

In particular, the scalar products $(\mathbf{e}_i, \mathbf{e}_j)$ of the base vectors $\mathbf{e}_i$ and $\mathbf{e}_j$ are equal to g_{ij} :

$$(\mathbf{e}_i, \mathbf{e}_j) = g_{ij} \quad (8.44)$$

(this follows from the equality $(\mathbf{e}_i, \mathbf{e}_j) = A(\mathbf{e}_i, \mathbf{e}_j)$ and from formula (8.42); by the way, formula (8.44) can be easily obtained directly from (8.43)).

Let us now consider a reciprocal basis $\mathbf{e}^1, \mathbf{e}^2, \dots, \mathbf{e}^n$ alongside the basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$. Let $\mathbf{x} = x_j \mathbf{e}^j$ and $\mathbf{y} = y_j \mathbf{e}^j$ be representations of the vectors $\mathbf{x}$ and $\mathbf{y}$ in terms of vectors of the reciprocal basis. Then we get the following formula for the scalar product $(\mathbf{x}, \mathbf{y})$:

$$(\mathbf{x}, \mathbf{y}) = A(\mathbf{x}, \mathbf{y}) = A(x_i \mathbf{e}^i, y_j \mathbf{e}^j) = A(\mathbf{e}^i, \mathbf{e}^j) x_i y_j.$$

We introduce the designation

$$g^{ij} = A(\mathbf{e}^i, \mathbf{e}^j) \quad (8.45)$$

and get the following expression for $(\mathbf{x}, \mathbf{y})$:

$$(\mathbf{x}, \mathbf{y}) = g^{ij} x_i y_j. \quad (8.46)$$

As in 8.2.2 (see Example 3), it is easy to verify that g^{ij} are the co-ordinates of a tensor of type (0, 2) which is symmetric with respect to the indices i and j . This tensor of type (0, 2) is also called a **metric tensor** of the space E^n .

We denote it by the same symbol G as the metric tensor of type (2, 0) introduced above: in 3.2 we shall find that the coordinates g_{ij} and g^{ij} can be regarded as the covariant and contravariant coordinates of the same vector. In what follows, we shall call the coordinates g_{ij} and g^{ij} the covariant and contravariant coordinates of the tensor G .

At the end of 8.1.2, while investigating the problem of constructing reciprocal bases, we introduced the quantities g_{ij} and g^{ij} by formulas (8.10). Comparing these formulas with (8.42) and (8.45), we infer that these quantities are the covariant and contravariant coordinates of the metric tensor G .

We have also proved in 8.1.2 that the matrices whose elements are the coordinates g_{ij} and g^{ij} , are mutually inverse. This means that there holds a relation

$$g^{ij} g_{jk} = \delta_k^i. \quad (8.47)$$

Thus the coordinates g^{ij} of the tensor G can be constructed from the coordinates g_{ij} and vice versa (for that purpose, we must use the familiar method of constructing the elements of an inverse matrix).

8.3.2. The operations of lifting and lowering indices with the aid of a metric tensor. The metric tensor G is used for the lifting and lowering of indices of the coordinates of a given tensor A .

The operation consists in the following.

Let A be a tensor of type (p, q) with coordinates $A_{i_1 i_2 \dots i_p}^{k_1 k_2 \dots k_q}$. We describe, for example, an operation of lifting the index i_1 . We contract the tensors G and A with respect to the superscript j in the first tensor and with respect to the subscript i_1 of the second tensor, i.e. construct a tensor with coordinates $g^{i\alpha} A_{\alpha i_2 \dots i_p}^{k_1 k_2 \dots k_q}$ and designate the index i in the coordinates of the tensor obtained as i_1 . Then we use the symbol $A_{i_2 \dots i_p}^{i_1 k_1 k_2 \dots k_q}$ for these coordinates. Thus we have

$$A_{i_2 \dots i_p}^{i_1 k_1 k_2 \dots k_q} = g^{i_1 \alpha} A_{\alpha i_2 \dots i_p}^{k_1 k_2 \dots k_q}. \quad (8.48)$$

Remark I. Since the sequence of indices in the coordinates of a tensor determines the “*numeration*” of its coordinates, it follows, in a general case, that when we lift an index, we must mark the place, among the superscripts, to which the given subscript will be lifted. It is sometimes necessary to mark the place of the index being lifted among the subscripts. This is done with the aid of a dot which is put into the place of the lifted index. Therefore, the coordinates of the tensor on the left-hand side of (8.48) should have been written as $A_{i_2 \dots i_p}^{i_1 k_1 k_2 \dots k_q}$.

For example, if we lift the *first* subscript to the second place among the superscripts, we get a tensor with the coordinates $A_{i_2 \dots i_q}^{k_1 i_1 k_2 \dots k_q}$.

Remark 2. The operation of lowering an index with the aid of a (fundamental tensor G can be defined by analogy. For instance, the coordinates of a tensor obtained by lowering the index k_q of the tensor A to the last place among the subscripts have the following form:

$$A_{i_1 \dots i_p k_q}^{k_1 \dots k_{q-1}} = g_{k_q \alpha} A_{i_1 \dots i_p}^{k_1 \dots k_{q-1} \alpha}.$$

Remark 3. The operation of lifting or lowering an index can be performed several times, each time for a different index of a given tensor.

Let us consider some examples of lifting and lowering indices of tensors.

Assume that $\mathbf{x}$ is a vector, and x_i and x^i are, respectively, its co-variant and contravariant coordinates (recall that a vector is a tensor of rank 1).

Let us lift the index i in the coordinates x_i by means of the fundamental tensor G .

As a result we get a tensor with the coordinates $g^{i\alpha} x_\alpha$. Since $x_\alpha = (\mathbf{x}, \mathbf{e}_\alpha)$ we have $g^{i\alpha} x_\alpha = g^{i\alpha} (\mathbf{x}, \mathbf{e}_\alpha) = (\mathbf{x}, g^{i\alpha} \mathbf{e}_\alpha)$.

According to (8.11), $g^{i\alpha} \mathbf{e}_\alpha = \mathbf{e}^i$ and $(\mathbf{x}, \mathbf{e}^i) = x^i$. Therefore, $g^{i\alpha} x_\alpha = x^i$.

Thus, the contravariant coordinates x^i of the vector $\mathbf{x}$ can be obtained as a result of an operation of lifting an index of the covariant coordinates x_i of that vector.

The covariant coordinates x_i can be obtained as a result of an operation of lowering an index of the contravariant coordinates x^i .

Let us find the result of an operation of lifting an index applied two times to the covariant coordinates g_{ij} of metric tensor G with the aid of the contravariant coordinates g^{ij} of that tensor. In other words, we shall find what the tensor with the coordinates

$$g^{i\alpha} g^{j\beta} g_{\alpha\beta} \quad (8.49)$$

is like.

Using the symmetry of the tensor G with respect to the subscripts and relation (8.47), we find that $g^{j\beta} g_{\alpha\beta} = g^{j\beta} g_{\beta\alpha} = \delta_{\alpha}^j$. Substituting the expression found for $g^{j\beta} g_{\alpha\beta}$ into (8.49) and using the properties of the Kronecker delta δ_{α}^j , we get

$$g^{i\alpha} g^{j\beta} = g^{ij}.$$

By a complete analogy we can verify the validity of the equality

$$g_{i\alpha} g_{j\beta} = g_{ij}.$$

The last two formulas emphasize once again that it is natural to regard g_{ij} and g^{ij} as covariant and contravariant coordinates of the metric tensor G .

8.3.3. Orthonormal bases in E^n . We have elucidated that the scalar product $(\mathbf{x}, \mathbf{y})$ in E^n can be defined by the metric tensor G whose co-ordinates g_{ij} are the elements of the symmetric positive definite matrix (g_{ij}) . Namely, according to (8.43),

$$(\mathbf{x}, \mathbf{y}) = g_{ij} x^i y^j.$$

It is known that by means of a transformation of the basis, the matrix of the bilinear form $g_{ij} x^i y^j$ can be reduced to a diagonal form. In this case, by virtue of the positive definiteness of the matrix, after the reduction of the matrix (g_{ij}) to a diagonal form, the coordinates of the metric tensor are equal to zero for $i \neq j$ and to unity for $i = j$. We symbolize these coordinates as g_{ij} and obtain

$$g_{ij} = \begin{cases} 0 & \text{for } i \neq j, \\ 1 & \text{for } i = j. \end{cases} \quad (8.50)$$

The basis $\mathbf{e}_i$ in which the coordinates g_{ij} of the metric tensor satisfy condition (8.50) is orthonormal. Indeed, since $(\mathbf{e}_i, \mathbf{e}_j) = g_{ij}$ (see (8.44)), we get, according to (8.50),

$$(\mathbf{e}_i, \mathbf{e}_j) = \begin{cases} 0 & \text{for } i \neq j, \\ 1 & \text{for } i = j, \end{cases}$$

and this means that $\mathbf{e}_i$ is an orthonormal basis.

We found in Chapter 4 that the scalar product $(\mathbf{x}, \mathbf{y})$ of the vectors $\mathbf{x}$ and $\mathbf{y}$ with the coordinates x^i and y^j could be calculated in an orthonormal basis by the formula

$$(\mathbf{x}, \mathbf{y}) = \sum_{i=1}^n x^i y^i, \quad (8.51)$$

and the square of the length $(\mathbf{x}, \mathbf{x})$ of the vector $\mathbf{x}$ by the formula

$$(\mathbf{x}, \mathbf{x}) = \sum_{i=1}^n (x^i)^2. \quad (8.52)$$

Let us direct our attention to the so-called **orthogonal linear transformations**, i.e. linear transformations under which an orthonormal basis passes into an orthonormal one. In other words, if L is an orthogonal -transformation and $\mathbf{e}_i$ is an orthonormal basis, the $L\mathbf{e}_i$ also forms an orthonormal basis.

Let us see how the transformation L acts on an arbitrary vector $\mathbf{x} = x^i \mathbf{e}_i$. We denote by $\mathbf{X}$ the result of the action of L on $\mathbf{x}$:

$$\mathbf{X} = L\mathbf{x}.$$

Using the property of linearity of L , we find that

$$\mathbf{X} = Lx^i \mathbf{e}_i = x^i L\mathbf{e}_i.$$

Since Le_i is a basis, it follows from the last relation that the vector X has the same coordinates in the basis Le_i as the vector x in the basis e_i , i.e. the coordinates of the vector retain their values under an orthogonal transformation.

Since Le_i is an orthonormal basis, the scalar product (X, Y) of the vectors $X = Lx$ and $Y = Ly$ can be found by formula (8.51), and the square of the length (X, X) of the vector $X = Lx$ by formula (8.52). We have found that the coordinates of a vector retain their value under an orthogonal transformation. From this and from relations (8.51) and (8.52) we obtain

$$(X, Y) = (X, y), (X, X) = (x, x).$$

Thus, the lengths of vectors and their scalar products do not change under an orthogonal transformation.

As is known, the orthogonal transformations L can be defined by means of an orthogonal matrix. The determinant $\det L$ of such a matrix satisfies the condition $\det L = \pm 1$.

Let us choose one of the orthonormal bases and agree to call it a *right* basis. In that case we say that the Euclidean space E^n is *oriented*. All the bases in E^n which are obtained from the given basis by means of orthogonal transformations with a determinant equal to $+1$ are called *right* bases and all bases which result from the given basis by means of orthogonal transformations with the determinant equal to -1 are called *left* bases. It is easy to verify when a right basis is transformed into a right one the determinant of the transformation is equal to $+1$ and when a left basis is transformed into a left one the determinant is equal to -1 .

We denote by $O(n)$ the set of all orthogonal transformations in E^n and by $O_+(n)$ the set of all orthogonal transformations of right bases.

We shall consider these sets in the next chapter.

Remark. In what follows, we call the arbitrary basis $e_1, e_2, \dots, e_n$ the *right basis* (the *left basis*) if the determinant of the transition matrix from the chosen orthonormal basis to the basis $e_1, e_2, \dots, e_n$ is positive (negative).

8.3.4. A discriminatory tensor. Let us consider the so-called **well-skew-symmetric tensor** $\varepsilon_{i_1 i_2 \dots i_p}$ of type $(p, 0)$, i.e. a tensor which is skew-symmetric with respect to any two subscripts.

For this tensor to be nonzero, it is necessary that the number p should not exceed n , i.e. satisfy the condition $p \leq n$ since in the case of $p > n$ any coordinate has at least two identical indices upon an interchange of which this coordinate must simultaneously change sign and remain unchanged. And this is only possible when this coordinate is equal to zero. Consequently, when $p > n$, any coordinate of a tensor is equal to zero, i.e. the tensor is zero.

Of special interest is a well-skew-symmetric tensor whose rank p is equal to the dimension n of the space.

Any coordinate $\varepsilon_{i_1 i_2 \dots i_p}$ in of such a tensor can be found by the formula

$$\varepsilon_{i_1 i_2 \dots i_p} = \begin{cases} 0, & \text{if at least two of the indices } i_1, i_2, \dots, i_p, \text{ coincide,} \\ (-1)^{\text{sign } \sigma} \varepsilon_{12 \dots n} & \text{if all the indices are different.} \end{cases} \quad (8.53)$$

In formula (8.53) sign σ is equal to 0 or +1 according as the interchange $\sigma = (i_1, i_2, \dots, i_p)$ is even or odd (sign σ is also called the *sign* of that interchange).

Let us consider some right orthonormal system of coordinates and set

$$\varepsilon_{12 \dots n} = 1 \quad (8.54)$$

in it. In this coordinate system relation (8.53) defines all the coordinates $\varepsilon_{i_1 i_2 \dots i_n}$ of a well-skew-symmetric tensor and, consequently, the tensor itself. In what follows we shall call this tensor a **discriminatory tensor**. We denote the coordinates of this tensor in the arbitrary basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ by $c_{i_1 i_2 \dots i_n}$.

We shall use the symbol B for the transition matrix from the chosen right orthonormal basis to some basis $\mathbf{e}_{1'}, \mathbf{e}_{2'}, \dots, \mathbf{e}_{n'}$ and the symbol b_i^j for the elements of that matrix.

In accordance with (8.53), to calculate the coordinates $c_{i'_1 i'_2 \dots i'_n}$ of the discriminatory tensor in the basis $\mathbf{e}_{1'}, \mathbf{e}_{2'}, \dots, \mathbf{e}_{n'}$ it is sufficient to know the value of the coordinate $c_{1' 2' \dots n'}$.

Using formula (8.19) for the transformation of the coordinates of a tensor and relation (8.54), and passing from the chosen orthonormal basis to the basis $\mathbf{e}_{1'}, \mathbf{e}_{2'}, \dots, \mathbf{e}_{n'}$ we get

$$\begin{aligned} c_{1' 2' \dots n'} &= b_1^{i_1} b_2^{i_2} \dots b_n^{i_n} \varepsilon_{i_1 i_2 \dots i_n} \\ &= \varepsilon_{12 \dots n} \sum_{\sigma = (i_1, i_2, \dots, i_n)} (-)^{\text{sign } \sigma} b_1^{i_1} b_2^{i_2} \dots b_n^{i_n} \\ &= \det (b_i^{j'}) = \det B. \end{aligned} \quad (8.55)$$

Assume that $g_{i'j'}$ are the coordinates of the fundamental tensor in basis $\mathbf{e}_{1'}, \mathbf{e}_{2'}, \dots, \mathbf{e}_{n'}$. Since $\tilde{G} = (g_{i'j'})$ is the matrix of the bilinear form $g_{i'j'} x^{i'} x^{j'}$, which is the scalar product of the vectors $\mathbf{x}$ and $\mathbf{y}$ with the coordinates $x^{i'}$ and $y^{j'}$ the formula $\tilde{G} = B' I B$ is valid when we pass from the given orthonormal basis (in which the matrix I of the bilinear form being considered is an identity matrix) to the basis $\mathbf{e}_{1'}, \mathbf{e}_{2'}, \dots, \mathbf{e}_{n'}$. Hence it follows that $\det \tilde{G} = \det B' I \det B = (\det B)^2$.

Designating $\det \tilde{G}$ as g , we get $\det B = \pm \sqrt{g}$ from the last relation. Using relations (8.55), we find that $c_{1' 2' \dots n'} = \pm \sqrt{g}$. Thus, in the arbitrary basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ the expression for the coordinate $c_{1 2 \dots n}$ of the discriminatory tensor has the form

$$c_{1 2 \dots n} = \pm \sqrt{g}, \quad (8.56)$$

where g is the determinant of the matrix (g_{ij}) of the metric tensor in the basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$.

Note that the plus sign in formula (8.56) corresponds to the right basis and the minus sign to the left basis.

8.3.5. An oriented volume. Let us introduce, in the oriented Euclidean space E^n , the so-called **affine system of coordinates** defining it as the set of a fixed point O with the coordinates $(0, 0, \dots, 0)$ and a basis $e_1, e_2, \dots, e_n$. The coordinates of any point M in E^n are defined in this case as the coordinates in the basis $e_1, e_2, \dots, e_n$

of the vector $\vec{OM}$.

Let us consider in E^n a numbered system of n vectors

$$\overset{1}{x}, \overset{2}{x}, \dots, \overset{n}{x} \quad (8.57)$$

and various vectors $\vec{OM}$ defined by the relations

$$\vec{OM} = \alpha_1 \overset{1}{x} + \alpha_2 \overset{2}{x} + \dots + \alpha_n \overset{n}{x}, \quad (8.58)$$

for various α_i satisfying the inequalities $0 \leq \alpha_i \leq 1, i = 1, 2, \dots, n$.

The set of all points M of the space E^n defined by relations (8.58) forms the so-called **n -dimensional parallelepiped** in E^n which is spanned by vectors (8.57).

The **oriented volume** $V(\overset{1}{x}, \overset{2}{x}, \dots, \overset{n}{x})$ of that parallelepiped is the number

$$V(\overset{1}{x}, \overset{2}{x}, \dots, \overset{n}{x}) = c_{i_1 i_2 \dots i_n} \overset{1}{x}^{i_1} \overset{2}{x}^{i_2} \dots \overset{n}{x}^{i_n}. \quad (8.59)$$

In this case, $c_{i_1 i_2 \dots i_n}$ are the coordinates of the discriminatory tensor in the basis

$e_1, e_2, \dots, e_n$, and $\overset{1}{x}^{i_1}, \overset{2}{x}^{i_2}, \dots, \overset{n}{x}^{i_n}$ are the contravariant coordinates of the vectors

$\overset{1}{x}, \overset{2}{x}, \dots, \overset{n}{x}$ in that basis.

We use the term “oriented volume” since in the case when vectors (8.57) form a right basis the oriented volume is positive ($V > 0$) and in the case of a left basis it is negative ($V < 0$).

Note that for $n = 3$ the oriented volume calculated for $n = 3$ by formula (8.59) is an ordinary volume of a parallelepiped spanned by the vectors $\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^n$. The volume is taken with the plus sign if the triad $\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^n$ is right-handed and with the minus sign if the triad is left-handed.

8.3.6. A vector product. We can use a discriminatory tensor to write a vector product in the three-dimensional space E^3 in tensor form. Such a notation is widely used in various calculations in the so-called curvilinear coordinates.

Suppose c_{ijk} are the coordinates of a discriminatory tensor in the given basis $\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3$ of the space E^3 . Let us lift the first index i of this tensor with the aid of the metric tensor g_{lm} , i.e. consider the tensor c^i_{jk} . Then the coordinates z^i of the vector $\mathbf{z} = [\mathbf{x} \mathbf{y}]$ (i.e. of the vector product of the vectors $\mathbf{x}$ and $\mathbf{y}$) in the basis $\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3$ have the form

$$z^i = c^i_{jk} x^j y^k. \quad (8.60)$$

Since $z^i = c^i_{jk} x^j y^k$ is a tensor of type (0, 1), z^i can be regarded as the contravariant coordinates of the vector. Therefore, to ascertain that z^i are indeed the coordinates of a vector product, it is sufficient to carry out the verification in some definite system of coordinates. This verification is elementary for an orthonormal basis and we invite the reader to carry it out.

Relation (8.60) can serve as the basis for the introduction of the vector product of $n - 1$ vectors in E^n .

Suppose $\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^{n-1}$ are some $n - 1$ vectors in E^n . Let us find the coordinates

z^i of the vector product $\mathbf{z} = \left[\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^{n-1} \right]$ by means of the relations

$$z^i = c_{i_1 i_2 \dots i_{n-1}}^i \mathbf{x}^{i_1} \mathbf{x}^{i_2} \dots \mathbf{x}^{i_{n-1}}. \quad (8.61)$$

In relations (8.61) $c_{i_1 i_2 \dots i_{n-1}}^i$ are the coordinates of a discriminatory tensor with

the lifted first index, and $\mathbf{x}^{i_1}, \mathbf{x}^{i_2}, \dots, \mathbf{x}^{i_{n-1}}$ are the contravariant coordinates of

the vectors $\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^{n-1}$.

8.3.7. A triple vector product. The reader is sure to be familiar,

from vector algebra, with the following formula for the triple vector

product $[a[bd]]$ of the vectors $\mathbf{a}$, $\mathbf{b}$, and $\mathbf{d}$:

$$[a[bd]] = \mathbf{b}(\mathbf{a}\mathbf{d}) - \mathbf{d}(\mathbf{a}\mathbf{b}). \quad (8.62)$$

Using relation (8.60) and formula (8.43) for the scalar product of vectors rewrite (8.62) as follows:

$$c_{kl}^i a^k c_{mn}^l b^m d^n = b^i g_{kl} a^k d^l - d^i g_{kl} a^k b^l. \quad (8.63)$$

With the aid of (8.63), we get a formula relating the tensors c_{kl}^i and g_{ik} , which we shall use for writing the coordinates of the triple vector product.

Let us perform the following transformations in formula (8.63). In the first term $b^i g_{kl} a^k d^l$ on the right-hand side of (8.63) we replace b^i by $b^m \delta_m^i$ (* The relation $b^i = b^m \delta_m^i$ follows from the properties of the Kronecker delta.) and replace the sum index l by n . In the second term on the right-hand side of (8.63) we set $d^i = b^m \delta_n^i$ and replace the sum index l by m . After these transformations, formula (8.63) assumes the form

$$(c_{kl}^i c_{mn}^l) a^k b^m d^n = (g_{kn} \delta_m^i - g_{km} \delta_n^i) a^k b^m d^n. \quad (8.64)$$

Since relation (8.64) is valid for any vectors $\mathbf{a}$, $\mathbf{b}$, and $\mathbf{d}$, it is an identity with respect to the coordinates a^k , b^m and d^n of these vectors and, therefore, the equality

$$c_{kl}^i c_{mn}^l = g_{kn} \delta_m^i - g_{km} \delta_n^i \quad (8.65)$$

holds for any indices i, k, m and n . We designate the coordinates of the triple vector product $[\mathbf{a}[\mathbf{b}\mathbf{d}]]$ as z^i . Then, according to (8.63), $z^i = c_{kl}^i c_{mn}^l a^k b^m d^n$. From this and from (8.65) we get the following expression for the coordinates z^i of the triple vector Product $[\mathbf{a}[\mathbf{b}\mathbf{d}]]$:

$$z^i = (g_{kn} \delta_m^i - g_{km} \delta_n^i) a^k b^m d^n. \quad (8.66)$$

Formula (8.66) is convenient for various applications.

8.4. A Metric Tensor of a Pseudo-Euclidean Space

8.4.1. The concept of a pseudo-Euclidean space and a metric tensor of a pseudo-Euclidean space. Let us consider an n -dimensional linear space L in which a non-singular symmetric bilinear form $A(\mathbf{x}, \mathbf{y})$ is defined which is the polar of an alternating quadratic form.

The **scalar product** $(\mathbf{x}, \mathbf{y})$ of the vectors $\mathbf{x}$ and $\mathbf{y}$ is the value $A(\mathbf{x}, \mathbf{y})$ of the bilinear form. The name “scalar product” is conventional since in this case the fourth axiom of the scalar product does not hold true. Namely, in the case when the bilinear form $A(\mathbf{x}, \mathbf{y})$ is the polar of the alternating quadratic form, the expression $A(\mathbf{x}, \mathbf{y})$ can have either a positive or a negative value depending on the choice of $\mathbf{x}$ and vanishes for nonzero values of $\mathbf{x}$. All the same we shall use the term “scalar product” since it is customary.

Here is the definition of a pseudo-Euclidean space.

Definition. A *pseudo-Euclidean space* is an n -dimensional linear space L in which a scalar product is defined by means of a non-singular symmetric bilinear form $A(\mathbf{x}, \mathbf{y})$ which is the polar of an alternating quadratic form.

The number n is the **dimension** of the pseudo-Euclidean space.

Let us isolate a basis $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ in the linear space L and denote by (g_{ij}) the matrix of the bilinear form $A(\mathbf{x}, \mathbf{y})$ in that basis (recall that $g_{ij} = A(\mathbf{e}_i, \mathbf{e}_j)$). If x^i and y^j are the contravariant coordinates of the vectors $\mathbf{x}$ and $\mathbf{y}$, then

$$A(\mathbf{x}, \mathbf{y}) = g_{ij} x^i y^j. \quad (8.67)$$

By a complete analogy with the arguments presented in 8.2.2 we can prove that g_{ij} are the coordinates of the tensor G of type $(2, 0)$. We shall call this tensor a **metric tensor of a pseudo-Euclidean space**.

Since the scalar product $(\mathbf{x}, \mathbf{y})$ is equal to $A(\mathbf{x}, \mathbf{y})$, we have, in accordance with (8.67), $(\mathbf{x}, \mathbf{y}) = g_{ij} x^i y^j$.

It is known that the matrix (g_{ij}) of the bilinear form $A(\mathbf{x}, \mathbf{y})$ can be reduced to a diagonal form. In that case, since the form $A(\mathbf{x}, \mathbf{y})$ is non-singular, the coordinates g_{ij} of the metric tensor are equal to zero for $i \neq j$ after the reduction to a diagonal form, and to unity or minus unity for $i = j$. The number p of positive and the number q of negative diagonal elements does not depend on the way of reducing the matrix to a diagonal form, and $p + q = n$ since the form $A(\mathbf{x}, \mathbf{y})$ is non-singular.

These arguments explain the designation $E_{(p, q)}^n$ for an n -dimensional pseudo-Euclidean space.

It is natural to pose a question concerning the measurement of the lengths of vectors in a pseudo-Euclidean space.

In a Euclidean space with a metric tensor g_{ij} the square of the length of the vector $\mathbf{x}$ with the coordinates x^i is assumed to be equal to $g_{ij} x^i x^j$. If we define the square of the length $s^2(\mathbf{x})$ of the vector $\mathbf{x}$ by means of the relation

$$s^2(\mathbf{x}) = g_{ij} x^i x^j, \quad (8.68)$$

then, evidently (since the form $A(\mathbf{x}, \mathbf{y})$ is alternating), there are nonzero vectors with a positive square of length, with a negative square of length and with a zero square of length. Therefore, to obtain only real numbers as a measure of the length of vectors, the relation

$$\sigma(\mathbf{x}) = (\text{sgn } s^2(\mathbf{x})) \sqrt{|s^2(\mathbf{x})|} \quad (8.69)$$

is usually taken as the length of a vector.

In what follows we shall use the following terminology borrowed from the special theory of relativity: we shall call the nonzero vector $\mathbf{x}$ **time-like** if $\sigma(\mathbf{x}) > 0$ for that vector, **space-like** if $\sigma(\mathbf{x}) < 0$, and isotropic if $\sigma(\mathbf{x}) = 0$.

The following statement holds true.

The set of termini of all time-like (space-like, isotropic) vectors whose origins coincide with the arbitrary fixed point M of the pseudo-Euclidean space, forms a cone.

For definiteness, we shall prove the statement for the case of time-like vectors. It is evidently sufficient to prove that if $\mathbf{x}$ is a time-like vector, then the vector $\lambda\mathbf{x}$ is also time-like for any real $\lambda \neq 0$.

Since the coordinates of the vector $\lambda\mathbf{x}$ are equal to λx^i , it follows, according to (8.68), that $s^2(\lambda\mathbf{x}) = \lambda^2 s^2(\mathbf{x})$, i.e. $\text{sgn } s^2(\lambda\mathbf{x}) = \text{sgn } s^2(\mathbf{x})$. It follows from this and from (8.69) that the vector $\lambda\mathbf{x}$ is time-like.

We have proved the statement for time-like vectors. Reasoning by analogy, we can make sure that the statement is also true for the case of space-like and isotropic vectors.

The cone of time-like vectors is often designated as T (from the word “time”) and the cone of space-like vectors as S (from the word “space”).

8.42. Galilean coordinates. The Lorentz transformations. In the theory of pseudo-Euclidean spaces an important role is played by the systems of coordinates in which the *square of the interval* (as it is customary to call the square of the length of the vector $s^2(\mathbf{x})$) *has the form*

$$s^2(\mathbf{x}) = \sum_{i=1}^p (x^i)^2 - \sum_{i=p+1}^n (x^i)^2. \quad (8.70)$$

According to the terminology borrowed from physics, *such systems of coordinates are called **Galilean***.

The transformations of coordinates which retain expression (8.70) for $s^2(\mathbf{x})$ are called the **Lorentz transformations**.

In 8.13.3 we shall discuss the Lorentz transformations of the space $E_{(1,3)}^4$, called the **Minkowski space**. This space is of interest for physics since it is the space of events of the special theory of relativity.

Note that in a Minkowski space the numeration of the coordinates of a vector usually begins with a zero. Thus, according to (8.70), the square $s^2(\mathbf{x})$ of the interval in the space $E_{(1,3)}^4$ is written as follows:

$$s^2(\mathbf{x}) = (x^0)^2 - (x^1)^2 - (x^2)^2 - (x^3)^2. \quad (8.71)$$

In physics, for the sake of convenience, the coordinate x^0 is identified with the expression ct , where c is the speed of light and t is time : x^1, x^2, x^3 are called spatial variables.

In the Minkowski space, the cone T of time-like vectors decomposes into two distinct connected open components T^+ (the cone of the future) and T^- (the cone of the past); the cone S of space-like vectors forms a connected set.

An explanation is due of connected components T^+ and T^- . To do that we give a physical interpretation of the vector $\mathbf{x}$ with the coordinates $x^i, i = 0, 1, 2, 3$ in the $E_{(1,3)}^4$: this vector is characterized by the quantity $\Delta t = x^0 / c$ and the vector $\Delta \mathbf{r} = \{x^1, x^2, x^3\}$. Thus, regarding $\mathbf{x}$ as a displacement in $E_{(1,3)}^4$, we can assume that this displacement is characterized by the displacement Δt in the time and the displacement $\Delta \mathbf{r}$ in space.

The time-like vectors $\mathbf{x}$ are defined by the condition $s(\mathbf{x}) > 0$. In this case, evidently $|\Delta \mathbf{r}|/|\Delta t|, c$. If $|\Delta t| < 0$, then for the displacement of $\mathbf{x}$ we get an inequality $0 < |\Delta \mathbf{r}|/|\Delta t| < c$. By definition such a displacement of $\mathbf{x}$ belongs to T^+ and can be regarded as a displacement of a particle “into the future”. If $|\Delta t| < 0$, then the displacement of $\mathbf{x}$ belongs to T^- and can be regarded as a displacement of a particle “into the past” (the motion of anti-particles is explained in this way in physics).

T^+ and T^- are evidently two connected open components of the cone T . The arguments presented explain their names, the cone of the future and the cone of the past.

8.4.3. The Lorentz transformations of the space $E_{(1,3)}^4$. Let us consider a Galilean system of coordinates with the basis $\mathbf{e}_1$ in the pseudo-Euclidean space $E_{(1,3)}^4$. In such a system of coordinates, the square of the interval $s^2(\mathbf{x})$ has the form (8.71), and the matrix (g_{ij}) of the metric tensor has the form

$$J = (g_{ij}) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \\ 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix}. \quad (8.72)$$

We pass to a new Galilean system of coordinates with the basis $\mathbf{e}_{i'}$ and find the conditions which must be satisfied by the coefficients $b_{i'}^i$ of the matrix B of transformation of base vectors. Since

$$\mathbf{e}_{i'} = b_{i'}^i \mathbf{e}_i, \quad (8.73)$$

and since the matrix $(g_{i'j'})$ of the metric tensor in the basis $\mathbf{e}_i$, also has the form (8.72), we can use formulas $g_{i'j'} = b_{i'}^i b_{j'}^j g_{ij}$ for the transformation of the coordinates of a metric tensor and get the following system of equations for determining the coefficients $b_{i'}^i$ of the matrix B of transformation of base vectors

(see (8.73)); in this case, the indices i and i' run through the values 0, 1, 2, 3 (see (8.71)) :

$$\left. \begin{aligned} (b_{0'}^0)^2 - \sum_{\alpha=1}^3 b_{0'}^\alpha b_{0'}^\alpha &= 1 \\ b_{0'}^0 b_{\beta'}^0 - \sum_{\alpha=1}^3 b_{0'}^\alpha b_{\beta'}^\alpha &= 0, \\ b_{\gamma'}^0 b_{\beta'}^0 - \sum_{\alpha=1}^3 b_{\gamma'}^\alpha b_{\beta'}^\alpha &= \begin{cases} -1 & \text{for } \gamma' = \beta' \\ 0 & \text{for } \gamma' \neq \beta', \gamma', \beta' = 1, 2, 3. \end{cases} \end{aligned} \right\} \quad (8.74)$$

We can write relations (8.74) in matrix form. For that purpose, we Consider, alongside the matrix B , a matrix B^* which can be obtained from B by changing the sign of the elements in the last three columns and a subsequent transposition. Relations (8.74) can evidently be written in the form

$$B^* B = J, \quad (8.75)$$

where the matrix J is defined by relation (8.72) (for the sake of comparison, recall that the matrix C of the orthogonal transformations of a Euclidean space satisfies the relation $C' C = I$, where I is an identity matrix).

Since $\det B^* = -\det B$, and $\det J = -1$, it follows from relation (8.75) that $\det B^* B = \det B^* \det B = -(\det B)^2 = -1$, i.e.

$$\det B = \pm 1. \quad (8.76)$$

We denote by L the collection of all. the Lorentz transformations of the Minkowski space. From these general transformations we isolate those which map each vector from T^+ into a vector also belonging to T^+ . The collection of such transformations is usually called the **Lorentz transformations of the space $E_{(1,3)}^4$** and is symbolized as $L_{\uparrow}$.

The general Lorentz transformations for which $\det B = +1$ form a class L_+ of the so-called **proper Lorentz transformations**.

The class L_- of **improper Lorentz transformations** is characterized

by the relation $\det B = -1$. The reflection about three spatial axes $x^{1'} = -x^1$, $x^{2'} = -x^2$ and $x^{3'} = -x^3$ can serve as an example of transformations of this kind.

Let B be the matrix of an arbitrary improper Lorentz transformation, and P , the matrix of the reflection just considered. The product of the arbitrary improper transformation and the reflection just considered is, evidently, a proper transformation with the matrix $B' = PB$.

Since $P \cdot P = I$, where I is an identity matrix, it follows that

$$B = P^2 B = P (PB) = PB'.$$

Thus, any improper Lorentz transformation is the product of some proper transformation with the matrix B' and a reflection with the matrix P .

The intersection of the sets $L_{\mp}$ and L_+ is symbolized as $L_{\mp+}$. Some group properties of the sets L , $L_{\mp+}$, L_+ and $L_{\mp-}$ will be discussed in the next chapter.

In conclusion we seek the transformations $L_{\mp}$ which do not change the coordinates x^2 and x^3 . It is clear that these are the transformations $L_{\mp}$ of a two-dimensional pseudo-Euclidean subspace with the coordinates x^0 and x^1 , in which the square of the interval can be calculated by the formula $(x^0)^2 - (x^1)^2$.

Let us write formulas (8.74) for the case in question. We get

$$\left. \begin{aligned} (b_0^0)^2 - (b_0^1)^2 &= 1, \\ b_0^0 b_1^0 - b_0^1 b_1^1 &= 0, \\ (b_1^0)^2 - (b_1^1)^2 &= -1. \end{aligned} \right\} \quad (8.77)$$

Setting $b_1^0 / b_0^0 = \beta$, we find, from (8.77), the following expressions for the coefficients b_i^j of the transformation matrix B of the basis vectors e_0, e_1, e_2, e_3 into the base vectors $e_{0'}, e_{1'}, e_{2'}, e_{3'}$.

$$b_{0'}^0 = \pm \frac{1}{\sqrt{1-\beta^2}}, b_{0'}^1 = \pm \frac{\beta}{\sqrt{1-\beta^2}},$$

$$b_{1'}^0 = \pm \frac{\beta}{\sqrt{1-\beta^2}}, b_{1'}^1 = \pm \frac{1}{\sqrt{1-\beta^2}}.$$

In these formulas the sign is chosen from the condition under which the Lorentz transformation belongs to the class L_+ . Without going into the details of calculations, we write the final formula for the transformation of coordinates:

$$x^{0'} = \frac{x^0 - \beta x^1}{\sqrt{1-\beta^2}}, x^{1'} = \frac{x^1 - \beta x^0}{\sqrt{1-\beta^2}}, x^{2'} = x^2, x^{3'} = x^3. \quad (8.78)$$

We set $x^0 = ct, x^1 = x, x^2 = y, x^3 = z, x^{0'} = ct', x^{1'} = x', x^{2'} = y', x^{3'} = z'$ in (8.78). Then we must rewrite formulas (8.8) as follows:

$$t' = \frac{t - \frac{\beta}{c}x}{\sqrt{1-\beta^2}}, x' = \frac{\beta ct + x}{\sqrt{1-\beta^2}}, y' = y, z' = z. \quad (8.79)$$

Let us now elucidate the physical meaning of the constant β .

We assume that the point P is *stationary in the coordinate system* (t', x', y', z') . This means that the time t' varies whereas the spatial coordinates x', y', z' of that point are constant. We shall investigate the behaviour of the point P with respect to the system (t, x, y, z) . Differentiating the last three equations (8.79) and taking into account that $dx' = dy' = dz' = 0$, we find that $0 = \frac{-\beta c dt + dx}{\sqrt{1-\beta^2}}, 0 = dy, 0 = dz$.

Therefore, $\frac{dx}{dt} = \beta c, \frac{dy}{dt} = 0, \frac{dz}{dt} = 0$.

Consequently, *every point P , which is stationary in the coordinate system (t', x', y', z') (and, consequently, the whole system of coordinates), moves relative to the system (t, x, y, z) with a constant velocity $v = \beta c$ in the direction of the Ox axis.*

Thus $\beta = v/c$, where v is the velocity of the movement of the system (t', x', y', z') relative to the system (x, y, z, t) . Note, that $0 < \beta < 1$ since $0 < v < c$.

Let us rewrite formulas (8.79) as follows:

$$t' = \frac{t - \frac{v}{c^2}x}{\sqrt{1 - \frac{v^2}{c^2}}}, x' = \frac{-vt + x}{\sqrt{1 - \frac{v^2}{c^2}}}, y' = y, z' = z. \quad (8.80)$$

Formulas (8.80) are the transition formulas from the inertial coordinate system (t, x, y, z) to another inertial system (t', x', y', z') . They are known as the **Lorents formulas**.

8.5. The Tensor of the Moment of Inertia

Let us consider a solid fixed at the point O . Suppose $\vec{r} = \vec{OM}$ is a radius vector of the point M of that solid and v is the velocity of the point M .

As is known, the angular momentum N is defined by the relation $N = \int_V [\mathbf{r}\mathbf{v}] dm$,

where V is the volume of the solid and $dm = \rho dV$ (ρ is the density of the solid).

Designating the contravariant coordinates of the vector N as N^i and using formula (8.60) for a vector product, we obtain

$$N^i = \int_V c_{kl}^i \mathbf{r}^k v^l dm \quad (8.81)$$

(recall that $c_{kl}^i = g^{is} c_{skl}$, where c_{skl} are the coordinates of the discriminant tensor in the given basis of the space E^3 , see 8.3.6)

By Euler's theorem, there is an instantaneous axis of rotation of a solid. Designating the vector of the instantaneous angular velocity as ω we get $v = [\omega \mathbf{r}]$. Using formula (8.60) for a vector product once again, we obtain

$$v^l = c_{pn}^l \omega^p \mathbf{r}^n. \quad (8.82)$$

Substituting the expression obtained for v^l into the right-hand side of (8.81) and taking into account that ω^p is independent of the integration variables, we get the following expression for N^i .

$$N^i = \int_V c_{kl}^i c_{pn}^l \mathbf{r}^k \mathbf{r}^n \omega^p dm = \omega^p \int_V c_{kl}^i c_{pn}^l \mathbf{r}^k \mathbf{r}^n dm = \omega^p \mathbf{J}_p^i. \quad (8.83)$$

The tensor

$$\mathbf{J}_p^i = \int_V c_{kl}^i c_{pn}^l \mathbf{r}^k \mathbf{r}^n dm \quad (8.84)$$

appearing on the right-hand side of (8.83) is known as the **tensor of the moment of inertia**.

Let us transform expression (8.84) for the tensor of the moment of inertia. For that purpose, we direct our attention to formula (8.65). According to this formula $c_{kl}^i c_{pn}^l = g_{kn} \delta_p^i - g_{kp} \delta_n^i$. Therefore

$$\mathbf{J}_p^i = \int_V (g_{kn} \delta_p^i - g_{kp} \delta_n^i) \mathbf{r}^n \mathbf{r}^k dm. \quad (8.85)$$

If we lower the index i in (8.85) with the aid of a metric tensor, we get a frequently used formula for the coordinates of the twice covariant tensor of the moment of inertia:

$$J_{ip} = \int_V (\mathbf{r}^2 g_{ip} - \mathbf{r}_i \mathbf{r}_p) dm. \quad (8.86)$$

The tensor of the moment of inertia is widely used in mechanics of a solid body. To give an example, we shall write the expression for the kinetic energy T with the aid of this tensor. We have

$$T = \frac{1}{2} \int_V v^2 dm = \frac{1}{2} \int_V g_{ij} v^i v^j dm.$$

But since $v^i = c_{pn}^i \omega^p \mathbf{r}^n$, the expression for T assumes the form

$$T = \frac{1}{2} \int_V g_{ij} c_{pn}^i c_{kl}^j \omega^p \mathbf{r}^n \omega^k \mathbf{r}^l dm = \frac{1}{2} \omega^p \omega^k \int_V g_{ij} c_{pn}^i c_{kl}^j \mathbf{r}^n \mathbf{r}^l dm.$$

From this, in accordance with (8.84), we find that $T = \frac{1}{2} \omega^p \omega^k J_{pk}$.

Chapter – 9

ELEMENTS OF THE GROUP THEORY

In this chapter we shall get acquainted with the main concepts of the group theory and study some applications of the theory.

The importance of the group theory is determined by its numerous applications in physics.

9.1. The Concept of a Group. The Main Properties of Groups

9.1.1. Composition laws. We say that a **law of composition** is defined in the set A if we are given a *mapping* T of ordered pairs of elements from A into the set A . In this case, the element c from A , which is put into correspondence with the elements a, b from A with the aid of the mapping T , is called the **composition** or **wreath product** of those, elements.

The composition c of the elements a and b is symbolized as aTb :

$$c = aTb.$$

Other notations are also used for the composition of the elements a and b of the set A . Most frequently used are the *additive* notation $c = a + b$ and the *multiplicative* notation $c = ab$.

In the case of an additive notation, the corresponding law of composition is called **addition** and in the case of a multiplicative notation it is called **multiplication**. The law of composition is said to be **associative** if for any elements a, b and c of the set A there holds a relation

$$aT(bTc) = (aTb)Tc.$$

The composition law is said to be **commutative** if the relation $aTb = bTa$ holds for any pair $a, b \in A$.

The element e of the set A is called **neutral** relative to the law T if the relation $aTe = a$ holds for any element a of the set A .

Ordinary addition and multiplication in the set of real numbers can serve as examples of composition laws. Both these laws are commutative. The neutral element for addition is zero, for multiplication it is unity.

9.1.2. The concept of a group. Some properties of groups. We give the following definition.

Definition 1. *A set A in which the law of composition T is defined is called a **group** G if that law is associative, there is a neutral element e relative to the law T , and every element a of the set A has an inverse a^{-1} , i.e. there is an inverse element such that $aTa^{-1} = e$.*

If we use the multiplicative notation of the composition of elements. Definition 1 assumes the following form.

Definition 2. *A set A of elements $a, b, c, \dots$, in which a law of composition, called multiplication, is defined such that it puts every pair of elements a, b of the set A into correspondence with a definite element $c = ab$ of the set, is called a **group** G if the law satisfies the following requirements:*

- 1°. $a(bc) = (ab)c$ (associativity).
- 2°. There is an element e of the set A such that $ae = a$ for any element a of that set (the existence of a neutral element).
- 3°. For any element a of the set A there is an inverse element a^{-1} such that $aa^{-1} = e$.

It is customary to call the neutral element E the **identity** element of the group G .

If the composition law T acting in the group G is *commutative*, then the group G is called **commutative** or **Abelian**. The additive notation of the composition of elements is most often used for Abelian groups. In that case the neutral element of the Abelian group is called a **zero**.

Let us consider some examples of groups. T

- (1) The set Z of integers forms an Abelian group with respect to summation.

Indeed, the operation of summation of integers is evidently a law of composition. It is clear that this law is associative and commutative. The integer zero is a neutral element (zero). The integer $-a$ is an inverse element of the integer a .

(2) The set of positive real numbers forms an Abelian group with respect to multiplication. This operation is a law of composition. This law is evidently associative and commutative. The neutral element is the real number unity. The inverse element of the number $a > 0$ is $1/a$.

(3) A linear space forms an Abelian group with respect to addition of elements. This operation is a law of composition. According to the axioms of a linear space this law is associative and commutative. The neutral element is a zero element of the space, the inverse element of the element x is $-x$.

(4) Let BCD be an equilateral triangle (see Fig. 9.1). Let us

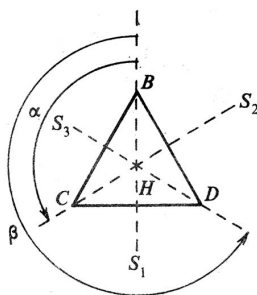


Fig. 9.1

consider the following set A of operations superimposing the triangle on itself

- 1°. The rotation α through $2\pi/3$ about the centre H mapping B in C .
- 2°. The rotation β through $4\pi/3$ mapping B into D .
- 3°. The symmetry S_1 mapping C into D .
- 4°. The symmetry S_2 mapping D into B .
- 5°. The symmetry S_3 mapping B into C .
- 6°. The identity operation 1.

The following table is a law of composition of the elements of the set A :

	1	α	β	S_1	S_2	S_3
1	1	α	β	S_1	S_2	S_3
α	α	β	1	S_2	S_3	S_1
β	β	1	α	S_3	S_1	S_2
S_1	S_1	S_3	S_2	1	β	α
S_2	S_2	S_1	S_3	α	1	β
S_3	S_3	S_2	S_1	β	α	1

(The rule for using the table can be easily seen from the successive performance of the operations S_1 and then $S_2 : S_2 S_1 = \beta$.)

This composition law is associative but not commutative, there is a neutral element which is an identity operation 1. Every operation has an inverse (there is an-identity operation in every row and every column of the table).

Thus the set A of the operations with the indicated law of composition is a group which is evidently noncommutative.

(5) **Permutation groups.**

A one-to-one mapping f of the arbitrary set E onto itself is called a permutation of the set E . In the process of permutation, every element a of the set E passes into the element $f(a)$, the inverse permutation f^{-1} maps $f(a)$ into a . The permutation $f(a) = a$ for any a of the set E is called an **identity permutation**.

If the set E consists of elements $a, b, c, \dots$, then the permutation f of that set is written as follows:

$$f = \begin{pmatrix} a & b & c & \dots \\ f(a) & f(b) & f(c) & \dots \end{pmatrix}.$$

The composition law is defined in a natural way in the set P of the permutations of the set E : if f_1 and f_2 are permutations of E , then the successive performance of $f_2 \circ f_1$ number of these permutations is a certain permutation of the set E . It is easy to see that the composition is associative.

If the set P contains an identity permutation, an inverse permutation of every permutation f and, in conjunction with any two permutations f_1, f_2 contains their composition $f_2 \circ f_1$, then, evidently, P is a group.

All the permutations of the set E form a group. For the finite set E consisting of n elements, this group is called a symmetric group S_n .

Let us return to example (4) in which we considered the operations of superimposition of the equilateral triangle BCD on itself. We designate the set of the vertices of the triangle as $E : E = (BCD)$.

The group of operations considered in Example (4) can evidently be obtained by the following group of permutations:

$$1 = \begin{pmatrix} B & C & D \\ B & C & D \end{pmatrix}, \quad \alpha = \begin{pmatrix} B & C & D \\ C & D & B \end{pmatrix}, \quad \beta = \begin{pmatrix} B & C & D \\ D & B & C \end{pmatrix},$$

$$S_1 = \begin{pmatrix} B & C & D \\ B & D & C \end{pmatrix}, \quad S_2 = \begin{pmatrix} B & C & D \\ D & C & B \end{pmatrix}, \quad S_3 = \begin{pmatrix} B & C & D \\ C & B & D \end{pmatrix}.$$

(6) Let us consider a group Z_2 consisting of two elements 0 and 1 in which multiplication is defined by the rule

$$0 \cdot 0 = 0, \quad 0 \cdot 1 = 1, \quad 1 \cdot 0 = 1, \quad 1 \cdot 1 = 0. \quad (9.1)$$

The unity of the group is the element 0.

This group is known as the **group of residue classes modulo 2**.

(7) Let us consider a group consisting of two elements: (a) an identity transformation of a Euclidean space (we denote this element by 0); (b) the mapping of a Euclidean space 'with respect to the origin' (we denote this element by 1).

The multiplication (i.e. the successive performance of operations (a) and (b)) of the elements 0 and 1 is evidently carried out by (rule (9.1)). We see that the group in question differs from the group properties of these two groups are the same.

Let us note the following *properties of groups* (we use the multiplicative notation for the composition).

Theorem 9.1. *If $aa^{-1} = e$, then $a^{-1}a = e$.*

Proof. Let x be an inverse of a^{-1} :

$$a^{-1}x = e.$$

Then $a = ae = a(a^{-1}x) = (aa^{-1})x = ex$, i.e. $a = ex$. Consequently, $a^{-1}a = a^{-1}(ex) = (a^{-1}e)x = e$, i.e. $a^{-1}a = e$. We have proved the theorem.

Theorem 9.2. *The relation $ea = a$ is valid for any element a of the group.*

Proof. By Theorem 9.1, $a^{-1}a = e$ and, in addition, $aa^{-1} = e$. Therefore, $ea = (aa^{-1})a = a = a$, i.e. $ea = a$. We have proved the theorem.

Theorem 9.3. *If $ax = e$ and $ay = e$, then $x = y$.*

Proof. Since $ay = e$, then y is an inverse of a and, therefore, according to Theorem 9.1, $ya = e$. Furthermore, we have $y = ye = y(ax) = (ya)x = ex = x$. We have proved the theorem.

The following important corollaries follow from the theorems we have proved.

Corollary 1. *The inverse of the element a^{-1} is an element a . Or, to put it otherwise, the element a^{-1} is the right as well as the left inverse element of the element a (i.e. $aa^{-1} = e$ and $a^{-1}a = e$).*

Corollary 2. *In any group the equations $ax = b$ and $ya = b$ can be solved uniquely. The solutions of these equations are the elements $x = a^{-1}b$ and $y = ba^{-1}$ respectively.*

Corollary 3. *There is a unique neutral element in a group (the unity of the group) (if $ae = a$ and $ae^* = a$, then $e = e^*$).*

Remark. *Note that the inverse element $(ab)^{-1}$ of the product ab is the element $b^{-1}a^{-1}$.*

Indeed, using the associative property of multiplication, we obtain $(ab)(b^{-1}a^{-1}) = a(bb^{-1})a^{-1} = aea^{-1} = aa^{-1} = e$.

9.1.3. Isomorphism of groups. Subgroups. The examples considered in 9.1.2 (see examples 4 and 5, examples 6 and 7) show that there are groups which differ by the nature of their elements but possess the same group properties. Such groups are naturally called **isomorphic**.

We give an exact definition of this concept.

Definition 1. Two groups G_1 and G_2 are said to be **isomorphic** if there exists a one-to-one mapping f of the group G_1 onto the group G_2 , such that the condition $f(ab) = f(a)f(b)$ holds for any elements a and b of G_1 .

Note that if e_1 is the unity of group G_1 , and e_2 is the unity of group G_2 , then $f(e_1) = e_2$. Indeed $f(e_1) = f(e_1 e_1) = f(e_1) \cdot f(e_1)$, and the multiplication by the inverse of $f(e_1)$ shows that $e_2 = f(e_1)$.

It should also be pointed out that the inverse mapping f^{-1} of group G_2 onto group G_1 for any elements x and y of G_2 satisfies the condition

$$f^{-1}(xy) = f^{-1}(x) f^{-1}(y).$$

In addition, it follows from the equation $e_2 = f(e_1)(aa^{-1}) = f(a)f(a^{-1})$ that for any a from G_1 the inverse of the element $f(a)$ is an element $f(a^{-1})$.

Thus, the isomorphic groups considered in abstract form, without indicating the nature of the elements are distinguished from the viewpoint of group properties.

Remark 1. The correspondence between the isomorphic groups G_1 and G_2 is called an **isomorphism** or **isomorphic mapping** of one group onto the other (the two groups being equivalent).

Remark 2. The isomorphic mapping of group G onto itself is called an **automorphism**.

The automorphisms of a group define in a definite way its symmetry.

If we consider separate automorphisms of a group as some elements and a successive performance of automorphism as the product of the corresponding

elements, then the automorphisms themselves form a group (whose unit element is an identity automorphism).

This group is known as the **group of automorphisms** of the given group.

It is easy to ascertain that ‘the group of automorphisms of the group Z_2 (see Example 6 in 9.1.2) is isomorphic to this group.

The concept of a **subgroup** plays an important part in the group theory.

Definition 2. *A subset G_1 of elements. of the group G is called a subgroup of that group if the following conditions are fulfilled: (1) if the elements a and b belong to G_1 , then ab also belongs to G_1 . (2) if the element a belongs to G_1 , then the inverse element a^{-1} also belongs to G_1 .*

A subgroup G_1 of the group G , regarded as a set in which the operation of multiplication is defined by the composition law from the enveloping group G , is a group.

It is easy to verify this assertion.

An elementary subgroup of any group is its unit element. The subgroup G_1 of all even numbers in the group G relative to the addition of all integers can serve as another example.

9.1.4. Cosets. Normal subgroups. Let H_1 and H_2 be arbitrary subsets of the group G .

The **product** of the subsets H_1 and H_2 is a subset H_3 consisting of all elements of the form $h_1 h_2$, where $h_1 \in H_1 \in H_2$.

The designation

$$H_3 = H_1 H_2$$

is used for the product of subsets.

Let us consider the case when H_1 consist of a single element h . Then, according to (9.2), the product of H_1 and H_2 can be written as hH_2 . Note that if the subsets

H_1 and H_2 are subgroups of the group G , then their product $H_1 H_2$ is not, in general, a subgroup.

Suppose H is a subgroup of the group G and a is an element of the group G . The set aH is said to be the **left coset** and the set Ha is said to be the **right coset** of the subgroup H in G .

Of course, when we choose another element instead of a , the right and left cosets of the subgroup H in G change a case.

Note the following properties of cosets (we formulate these properties only for the left cosets, the formulation for the right cosets is similar) :

- 1°. If $a \in H$, then $aH \equiv H$.
- 2°. The cosets aH and bH coincide if $a^{-1}b \in H$.
- 3°. Two cosets of the same subgroup H either coincide or do not possess any elements in common.
- 4°. If aH is a coset, then $a \in aH$.

The first property is obvious. Let us verify the validity of property 2°. Since $a^{-1}b \in H$ according to 1°, we have $bH = (aa^{-1})bH = a(a^{-1}b)H = aH$ because $aa^{-1} = e$. We have thus established property 2°.

Let us prove property 3°. It is, evidently, sufficient to prove that when the cosets aH and bH have an element in common they coincide.

Suppose the elements $h_1 \in H$ and $h_2 \in H$ are such that

$$ah_1 = bh_2 \quad (9.3)$$

(equation (9.3) means that the cosets aH and bH have an element in common). Since H is a subgroup of the group G , the element $h_1h_2^{-1}$ belongs to H . From this and from (9.3) we find that

$$a^{-1}b = h_1h_2^{-1} \in H.$$

Consequently, according to property 2°, $aH = bH$. We have proved property 3°.

Property 4° follows from the fact that the subgroup H contains a unit element e and, therefore, $ae = a \in aH$.

Let H be a subgroup of G for which all the left cosets are simultaneously the right cosets. In this case the relation

$$aH = Ha \quad (9.4)$$

must hold for any element a .

Indeed, according to property 4°, the element $a \in aH$. On the other hand, the class aH is simultaneously a certain Class Hb which evidently (since $a \in Hb$) coincides with the set Ha .

A subgroup H for which all the left-cosets are the right-cosets is called a **normal** or **invariant subgroup of group G** .

The following statement holds true.

If H is a normal subgroup of group G , then the product of the cosets is also a coset.

Indeed, let aH and bH be cosets. Then, by the definition of the product of cosets as subsets of the group G , with (9.4) taken into account, we get

$$aHbH = a (Hb) H = a (bH) H = (ab) (HH) = (ab) H.$$

i.e. the product of the cosets $aHbH$ is a coset $(ab) H$.

9.1.5. Homomorphisms. Factor groups.

Suppose G is a group with elements $a, b, c, \dots$, and $\bar{G}$ is a certain set in which the composition law of its elements $\bar{a}, \bar{b}, \bar{c}, \dots$ is defined. We shall use the multiplicative notation for the composition : $\bar{a} = \overline{ab}$, and the element c is the product of the elements $\bar{a}$ and $\bar{b}$.

Definition I. *The mapping f of the group G into the set $\bar{G}$*

$$f : G \rightarrow \bar{G} \quad (9.5)$$

*is called a **homomorphic mapping** if the relation*

$$f(ab) : f(a)f(b), \quad (9.6)$$

where $f(a)$, $f(b)$ and $f(ab)$ are the images of the elements a , b and ab under the mapping f , holds for any elements $a \in G$ and $b \in G$.

In this case $\bar{G}$ is a **homomorphic image** or **map** of G .

When $\bar{G}$ is a subset of G , the name **endomorphism** or **endomorphism mapping** is used for homomorphism (9.5).

Remark. If the homomorphic mapping (homomorphism) of the group G onto the set $\bar{G}$ is defined, then all the elements of the group are divided into nonintersecting classes: one class contains all the elements of G which are mapped into the same element of the set $\bar{G}$.

The following statement holds true. '

Theorem 9.4. *The homomorphic image of a group is a group.*

Proof. Suppose $\bar{a}, \bar{b}, \bar{c}, \dots$ are elements of the homomorphic image $\bar{G}$ of the group G under the homomorphic mapping f . This means that there are elements $a, b, c, \dots$ in the group G such that $\bar{a} = f(a)$, $\bar{b} = f(b)$, $\bar{c} = f(c)$, $\dots$

Then multiplication of elements in the set $\bar{G}$ agrees with rule (9.6).

Let us verify that this operation of multiplication meets requirements 1°, 2° and 3° of Definition 2 (see 9.1.2).

1°. *Associativity of multiplication.* Let us products $\bar{a}(\bar{b}\bar{c})$ and $(\bar{a}\bar{b})\bar{c}$. We have, according to (9.6),

$$\bar{a}(\bar{a}\bar{c}) = f(a)(f(b)f(c)) = f(a)f(bc) = f(abc),$$

$$(\bar{a}\bar{b})\bar{c} = (f(a)f(b))f(c) = f(ab)f(c) = f(abc).$$

Comparing these relations, we find $\bar{a}(\bar{b}\bar{c}) = (\bar{a}\bar{b})\bar{c}$. Consequently, the associativity of multiplication of the elements holds true.

2°. *The existence of unity.* We designate as $\bar{e}$ the element $f(e)$, where e is the unity of the group G : _

$$\bar{e} = f(e).$$

According to rule (9.6) we have $\bar{a} \bar{e} = f(a) f(e) = f(ae) = f(a) = \bar{a}$ for any element $\bar{a}$ of the set $\bar{G}$.

Consequently, the element $\bar{e}$ indeed plays the part of a unity.

3°. *The existence of an inverse element.* Let us symbolize as $\bar{a}^{-1}$ the element $f(a^{-1})$, where a^{-1} is the inverse of the element a in the group G .

According to (9.6) we have $\bar{a} \bar{a}^{-1} = f(a) f(a^{-1}) = f(aa^{-1}) = f(e) = \bar{e}$.

Consequently, the element $\bar{a}^{-1}$ plays the part of an inverse of the element $\bar{a}$.

Thus requirements 1°, 2°, 3° of Definition 2 of a group are met for the operation of multiplication of the elements of $\bar{G}$. Therefore, $\bar{G}$ is a group. We have proved the theorem.

Let H be a normal subgroup of the group G . Let us define the following mapping f of the group G onto the set $\bar{G}$ of cosets by the normal subgroup H : if a belongs to G , then we associate this element with the coset to which the indicated element belongs. According to property 3° of cosets (see 9.1.4), under such a mapping every element of the group G is associated with only one class, i.e. we indeed have a mapping f of the group G onto the set of cosets with respect to the normal subgroup H .

Let us prove the following theorem.

Theorem 9.5. *The mapping f of the group G onto cosets by the normal subgroup H indicated above is a homomorphic mapping when the multiplication of cosets is defined as that of subsets of*

the group G .

Proof. At the end of 9.1.4 we proved the statement that if aH and bH are cosets, then the product $aHbH$ of those cosets, as subsets of G , is a coset $(ab)H$. Consequently, by means of the mapping f being considered the product ab is put into correspondence with a coset $(ab)H$ equal to the product of the cosets aH and bH .

Therefore, f is a homomorphic mapping. We have proved the theorem.

Corollary. *The set of the group G by the normal subgroup H with the operation of multiplication of these cosets as subsets of G forms a group.*

This group is called a **factor group** of the group G by the normal subgroup H and is symbolized as G/H .

The validity of the corollary follows from Theorem 9.4.

Remark. The mapping f of the group G onto the set of cosets by the normal subgroup H is evidently a homomorphic mapping of that group onto the factor group G/H .

Let us consider the following example.

Let R^n be an n -dimensional linear coordinate space which, as was mentioned in Example 3 of 9.1.2, is an Abelian (i.e. commutative) group with respect to addition of elements (recall that the points x of that space are ordered collections of n real numbers $(x_1, \dots, x_n)$, the addition of elements $(x_1, \dots, x_n)$ and $(y_1, \dots, y_n)$ being carried out by the rule $(x_1 + y_1, \dots, x_n + y_n)$).

By definition, R^n is a direct product of one-dimensional spaces:

$$R^n = R_{(1)}^1 \times R_{(2)}^1 \times \dots \times R_{(n)}^1.$$

Since $R_{(n)}^1$, for example, is an Abelian subgroup, then, evidently, $R_{(n)}^1$ is a normal subgroup of the group R^n . The coset of the element a from R^n is a straight line passing through the point a parallel to the straight line $R_{(n)}^1$, and the factor group $R^n / R_{(n)}^1$ is isomorphic to the $(n - 1)$ dimensional subspace R^{n-1} :

$$R^{n-1} = R_{(1)}^1 \times R_{(2)}^1 \times \dots \times R_{(n-1)}^1. \quad (9.7)$$

Note that the designation of the factor group $R^n / R_{(n)}^1$ is explained in a certain way by means of the relation

$$R^n / R_{(n)}^1 = R_{(1)}^1 \times \dots \times R_{(n)}^1 / R_{(n)}^1 = R_{(1)}^1 R_{(2)}^1 \times \dots \times R_{(n-1)}^1, \quad (9-8)$$

which follows from relation (9.7). It should be pointed out that the last equality sign in formula (9.8) must be regarded as an isomorphism between the respective groups.

We have proved that the homomorphism of the group G onto the factor group G/H is defined by the normal subgroup H . The reverse statement is also true : if the homomorphism of the group G onto the set $\bar{G}$ is defined, then that homomorphism can be used to define the normal subgroup H such that $\bar{G}^*$ and the factor group G/H are isomorphic.

Let us prove two theorems related to this statement.

* According to Theorem 9.4, the homomorphic image of a group is a group.

Theorem 9.6. *Assume that f is a homomorphic mapping of the group G onto $\bar{G}$ and H is a set of elements of the group G which, under the homomorphic mapping f is mapped into the element $f(e)$, where e is the unity of the group G . Then, H is a normal subgroup of the group G .*

Proof. It is sufficient to prove that H is a subgroup of the group G and every left coset by that subgroup is simultaneously a right coset.

Let us first verify that H is a subgroup of the group G . For that purpose, we must prove that if $a \in H$ and $b \in H$, then $ab \in H$ and also if $a \in H$, then $a^{-1} \in H$.

Assume $a \in H$ and $b \in H$. Since f is a homomorphism we have $f(ab) = f(a) f(b) = f(e) f(e)$.

But $f(e)$ plays the part of a unity in the group $\bar{G}$ (see Theorem 9.4). Therefore, $f(e) f(e) = f(e)$, i.e. $f(ab) = f(e)$. Consequently, $ab \in H$.

Furthermore, assume that $a \in H$, i.e. $f(a) = f(e)$. Then, it (a^{-1}) is an inverse element of a , then $aa^{-1} = e$, i.e. $aa^{-1} \in H$. Since f is a homomorphism, we have $f(e) = f(aa^{-1}) = f(a) f(a^{-1}) = f(e) f(a^{-1}) = f(a^{-1})$. Therefore, $f(a^{-1}) = f(e)$ and, consequently, $a^{-1} \in H$.

Let us prove now that every left coset is at the same time a right coset.

Let a be an arbitrary element of the group G . We shall prove that the set A of the elements of the group G , which are mapped into the element $f(a)$ under the homomorphism f , is simultaneously the left aH and the right Ha coset. And this completes the proof of the theorem.

Assume $a' \in A$. Consider the equation*

$$ax = a'. \quad (9.9)$$

Since f is a homomorphism and $f(a') = f(a)$, we get from this equation $f(ax) = f(a)f(x) = f(a') = f(a)$, i.e. $f(x) = f(e)$. Therefore, $x \in H$. But then, according to (9.9), $a' = xa$, i.e. $a' \in Ha$.

Directing our attention to the equation $xa = a'$ and reasoning by analogy, we can verify that $x \in H$. But then $a' = xa$, i.e. $a' \in Ha$. Thus $A = aH = Ha$. We have proved the theorem.

Theorem 9.7 (theorem on the homomorphisms of groups). *Assume that f is a homomorphism of the group G onto G and H is the normal subgroup of the group G whose elements correspond to the unity of the group $\bar{G}^*$ under the homomorphism f . Then the group $\bar{G}$ and the factor group G/H are isomorphic.*

Proof. Let us establish a one-to-one correspondence between the elements of the group G and the cosets by the normal subgroup H : we put the element a of the group $\bar{G}$ into correspondence with the coset which is mapped into $\bar{a}$ by means of f . This correspondence is evidently one-to-one since, in accordance with property 3° of cosets (see 9.2.4), these cosets do not intersect. If we define the multiplication of these cosets as that of the subgroups of the group $\bar{G}$ and use the assertion proved at the end of 9.1.4, we shall easily see that the one-to-one correspondence we have just established is an isomorphism. But the cosets are exactly elements of a factor group. We have proved the theorem.

9.2. Transformation Groups

We shall study here groups of non-singular linear transformations of a linear and, in particular, Euclidean space.

*By virtue of Corollary 2 of Theorem 9.3, this equation is solvable. An element $x = a^{-1} a'$ is its solution.

9.2.1. Non-singular linear transformations. The concept of a linear operator was introduced in 5.1.1. Recall that a linear operator A is a mapping of a linear space V into a linear space W under which the image of the sum of elements is equal to the sum of their images and the image of the product of the elements by a number is equal to the product of that number by the image of the element.

We shall consider the so-called **non-singular linear operators** mapping the given finite-dimensional linear space V into the same space. The linear operator A is called *non-singular* if $\det A \neq 0$ *

Note the following *important property of non-singular operators: every such operator maps the space V onto itself one-to-one.*

In other words, *if A is a non-singular operator, then every element $x \in V$ is associated with only one element $y \in V$ which can be found by the formula*

$$y = Ax, \quad (9.10)$$

and if y is any fixed element of the space V , then there is only one element x such that $y = Ax$.

To prove the second part of the formulated assertion, we shall consider the matrix notation of the action of a linear operator. Thus, if $A = (a_k^j)$ is the matrix of the operator A in a given basis, and the elements x and y have the coordinates $x^1, \dots, x^n$ and $y^1, \dots, y^n$ respectively, then, according to formula (5.14) (see 5.2.1), relation (9.10) can be rewritten as

$$y^j = \sum_{k=1}^n a_k^j x^k, \quad j = 1, 2, \dots, n, \quad (9.11)$$

and, therefore, the coordinates x^k can be regarded as the unknowns with the specified coordinates y^j . Since the operator A is non-singular, i.e. $\det A \neq 0$, system of equations (9.11) has a unique solution for unknowns x^k . And this means that for every fixed element $y \in V$ there is a unique element x such that $y = Ax$.

*Recall that $\det \mathbf{A}$ was introduced in 5.2.2 as the determinant of the matrix of a linear operator in a given basis. It has also been proved there that the value of $\det \mathbf{A}$ does not depend on the choice of the basis.

Thus, the *result of the action of a non-singular linear operator can be regarded as a mapping of the linear space V onto itself.*

Therefore, *when a non-singular operator is specified, we can speak of a non-singular linear transformation of the space V or, simply, of a linear transformation of the space V .*

9.2.2. A group of linear transformations. Assume that V is an n -dimensional linear space with the elements $\mathbf{x}, \mathbf{y}, \mathbf{z}, \dots$ and $GL(n)$ is a set of all non-singular linear transformations of that space.

Let us determine in $GL(n)$ a composition law which we shall call **multiplication**. We define the multiplication of the linear transformations from $GL(n)$ in the same way as we defined the multiplication of linear operators in 5.1.2.

Namely, *the product $\mathbf{AB}$ on the linear transformations $\mathbf{A}$ and $\mathbf{B}$ belonging to the set $GL(n)$ is a linear operator acting according to the rule*

$$(\mathbf{AB})\mathbf{x} = \mathbf{A}(\mathbf{Bx}). \quad (9.12)$$

Note that in a general case $\mathbf{AB} \neq \mathbf{BA}$.

For the indicated product to be a composition law indeed (see 9.1.1), it is sufficient to prove. that the transformation $\mathbf{AB}$ is non-singular and this follows from the fact that the matrix of the linear operator $\mathbf{AB}$ is equal to the product of matrices of the transformations $\mathbf{A}$ and $\mathbf{B}$ and, consequently, $\det(\mathbf{AB}) = \det \mathbf{A} \cdot \det \mathbf{B} \neq 0$ since $\det \mathbf{A} \neq 0$ and $\det \mathbf{B} \neq 0$.

Let us prove now the following theorem.

Theorem 9.8. *The set $GL(n)$ of non-singular linear transformations of the linear n -dimensional space V with the operation of multiplication introduced above is a group (called a **group of linear transformations of the linear space V**).*

Proof. Let us verify requirements 1°, 2°, 3° of Definition 2 of a group (see 9.1.2).

1°. *Associativity of multiplication, i.e. the equation*

$$A (BC) = (AB) C$$

is valid since, according to (9.12), the product of linear transformations consists in their successive actions and, therefore, the linear transformations $A (BC)$ and $(AB) C$ coincide with the linear transformation ABC and are consequently identical.

2°. *The existence of a unity.* We denote an identity transformation by I . This transformation is non-singular since $\det I = 1$. The equality $AI = IA = A$ is evidently valid for any transformation A from $GL(n)$.

Consequently, the linear transformation I plays the part of a unity.

3°. *The existence of an inverse element.* Assume that A is any fixed non-singular linear transformation. Let us direct our attention to coordinate notation (9.11) of that transformation. Since $\det A \neq 0$, we can use system (9.11) to define uniquely the determinant x (coordinates x^k) from the given y (from the given coordinates y^j).

Consequently, we have defined the inverse transformation A^{-1} which is evidently linear (this follows from (9.11)); in addition $A^{-1} A = I$ by the definition itself. Therefore, the linear operator A^{-1} plays the part of the inverse operator for A .

Thus requirements 1°, 2°, 3° from Definition 2 of a group are met for the operation of multiplication of elements from $GL(n)$. Therefore, $GL(n)$ is a group. We have proved the theorem.

9.2.3. Convergence of elements in the group $GL(n)$. The subgroups of the group $GL(n)$. In this and the following subsections we shall consider the group $GL(n)$ in an n -dimensional Euclidean space V .

We introduce the concept of convergence in the group $GL(n)$.

Definition. *The sequence of the elements $\{A_n\}$ from $GL(n)$ is said to be convergent to the element $A \in GL(n)$ if the sequence $\{A_n x\}$ converges to Ax for any x^* .*

We shall use the concept of convergence in $GL(n)$ later on when we introduce the so-called *compact* groups.

The subgroups of $GL(n)$ are classified as follows.

1°. *Finite subgroups*, i.e. subgroups containing a finite number of elements.

A subgroup of reflections about the origin containing two elements, an identity transformation and a reflection about the origin (see Example 7 in 9.1.2) can serve as an example of a finite subgroup.

2°. *Discrete subgroups*, i.e. subgroups containing denumerably many elements.

A subgroup of rotations of a plane about the origin through the angles $k\varphi$, $k = 0, \pm 1, \pm 2, \dots$, where φ is an angle incommensurable with π , can serve as an example of such a subgroup.

3°. *Continuous subgroups*, i.e. subgroups containing more than a denumerable number of elements.

The subgroup of all rotations of a three-dimensional space about a fixed axis is an example of a continuous subgroup.

From the continuous subgroups of the group $GL(n)$ we can isolate the so-called compact subgroups, i.e. subgroups in which we can isolate, from any infinite set of elements, a sequence converging to an element of that subgroup.

9.2.4. A group of orthogonal transformations. In the group $GL(n)$ we can isolate a special subgroup of the so-called **orthogonal transformations**. These transformations, regarded as a separate set, form a group called an orthogonal group.

We introduce the concept of orthogonal transformations.

Recall that we consider non-singular linear transformations. The concept of such a transformation is equivalent to the notion of a non-singular operator, i.e. the operator A for which $\det A \neq 0$.

Recall now the concept of an orthogonal operator acting in the real Euclidean space V introduced in 5.9.

Namely, we called the linear operator P orthogonal if the relation

$$(Px, Py) = (x, y) \quad (9.13)$$

held for any x and y .

The result of the action of the orthogonal operator P is an orthogonal transformation P .

It was proved in Theorem 5.36 that the operator P was orthogonal if and only if there was an inverse operator P^{-1} and there held an equality

$$P^{-1} = P^* \quad (9.14)$$

In this equation P^* is an adjoint operator of P .

Thus, if the transformation P is orthogonal, then it has an inverse P^{-1} . Hence it follows that *every orthogonal transformation is non-singular*. Indeed, since $PP^{-1} = I$, where I is an identity transformation, we have

$$\det P \cdot \det P^{-1} = \det I = 1,$$

i.e. $\det P \neq 0$. Consequently, the orthogonal transformation P is non-singular.

Note the following important property of orthogonal transformations.

Theorem 9.9. *The set of all orthogonal transformations of the Euclidean space V with an ordinary operation of multiplication of linear transformations forms a group (called an **orthogonal** group and symbolized as $O(n)$).*

Proof. It is sufficient to prove that the product of orthogonal transformations is an orthogonal transformation. The existence of an inverse transformation (an inverse element) for a given orthogonal transformation was proved in Theorem 5.36.

Thus, let P_1 and P_2 be orthogonal transformations. Let us consider the product P_1P_2 . According to Theorem 5.36, we have only to prove the relation

$$(P_1P_2)(P_1P_2)^* = I. \quad (9.15)$$

We established in 5.5.1 (see property 5° of adjoint operators) that $(P_1P_2)^* = P_2^*P_1^*$.

Using this relation and the orthogonality of the transformations P_1 and P_2 , we get

$$(P_1P_2)(P_1P_2)^* = (P_1P_2)(P_2^*P_1^*) = P_1(P_2P_2^*)P_1^* = P_1IP_1^* = P_1P_1^* = I.$$

We have proved relation (9.15) and thus proved the theorem.

Remark 1. The orthogonal group is evidently a subgroup of the group $GL(n)$.

Remark 2. The value of the determinant $\det \mathbf{P}$ of the orthogonal transformation $\mathbf{P}$ satisfies the relation

$$(\det \mathbf{P})^2 = 1. \quad (9.16)$$

Thus,

$$\det \mathbf{P} = \pm 1. \quad (9.17)$$

To prove (9.16), we must keep in mind that the relation

$$\mathbf{P}\mathbf{P}' = \mathbf{I}, \quad (9.18)$$

where $\mathbf{P}'$ is a transposed matrix obtained from $\mathbf{P}$ by an interchange of rows and columns and $\mathbf{I}$ is an identity matrix, holds for the transformation $\mathbf{P}$.

Since $\det \mathbf{P} = \det \mathbf{P}'$ (a determinant does not change upon an interchange of rows and columns) and $\det \mathbf{I} = 1$, it follows from relation (9.18) that $(\det \mathbf{P})^2 = 1$, i.e. $\det \mathbf{P} = \pm 1$. Since by definition $\det \mathbf{P}$ is introduced as the determinant of the matrix $\mathbf{P}$ in any basis, we have proved relations (9.16) and (9.17).

Relation (9.17) for the determinant of an orthogonal transformation serves as the basis for the division of all these transformations into two classes.

In the first class we place all orthogonal transformations for which $\det \mathbf{P} = +1$. In what follows, we shall call these transformations **proper**.

In the second class we place all orthogonal transformations for which $\det \mathbf{P} = -1$. These transformations are called **improper**.

The set of all proper orthogonal transformations forms a group called a **special** or **proper orthogonal group**. It is symbolized as $SO(n)$.

We can prove that every group $SO(n)$ is compact.

9.2.5. Certain discrete and finite subgroups of an orthogonal group. We shall not try to make the exposition here as complete as possible. We shall cite examples showing the characteristics of some subgroups of an orthogonal group.

Note that finite and discrete subgroups of the group $O(3)$ find important application in crystallography.

1°. Let us consider a two-dimensional orthogonal group $O(2)$. From this group we can isolate a discrete group of rotations through the angle $k\varphi$, $k = 0, \pm 1, \pm 2, \dots$

We designate as a the element of this subgroup corresponding to the value $k = \pm 1$. Then the element a_k , corresponding to the rotation through the angle $k\varphi$ for $k > 0$, is evidently equal to $\underbrace{a \cdot a \dots a}_{k \text{ times}}$. This relation can be abbreviated as

$$a_k = a^k, k = 1, 2, \dots$$

If we symbolize as a^{-1} the inverse of the element a (a^{-1} is the element corresponding to the rotation through the angle $-\varphi$) and designate the unity of the subgroup being considered as a^0 , then, evidently, any element a_k for a positive, negative and zero value of k can be written as

$$a_k = a^k, k = 0, \pm 1, \pm 2, \dots \quad (9.19)$$

The groups whose elements a_k can be represented in the form (9.19) are called **cyclic groups**.

Cyclic groups are evidently discrete.

Note the following two types of the cyclic subgroups of rotation.

- (1) If $\varphi \neq 2\pi p / q$, where p and q are integers (i.e. the angle is incommensurable with π), then all the elements a_k are distinct.
- (2) If $\varphi = 2\pi p / q$, where p and q are coprime numbers, then the relation $a_{k+q} = a_k$ holds true, i.e. $a^q = a^0$.

The groups for which the last relation holds true are called the **cyclic groups of order q** . –

2°. Let us direct our attention now to the so-called subgroups of **mirror symmetry**.

Every subgroup of mirror symmetry consists of two elements: a unity (an identity transformation) and a reflection about some plane or about the origin. It is relatively easy to verify that an identity transformation and a reflection form a group. To do that, it is sufficient to note that two successive reflections produce an identity transformation (see Example 7 in 9.1.2).

Let us consider, for instance, a subgroup $\{I, P\}$ of the group $O(3)$ consisting of a unity I and a reflection P of a three-dimensional space about the origin. In an orthonormal basis the matrix P of that transformation has the form

$$P = \begin{pmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & -1 \end{pmatrix}.$$

Since the determinant $\det P = -1$, the subgroup $\{I, P\}$ is improper. We pointed out in Example 7 in 9.1.2 that the subgroup $\{I, P\}$ was isomorphic to the group Z_2 of residues modulo 2.

Let us prove the following statement.

The subgroup $\{I, P\}$ being considered is a normal subgroup of the group $O(3)$.

We have to prove that the relations

$$aI = Ia, aP = Pa \quad (9.20)$$

hold for any element a from $O(3)$ (these relations show that the left and right cosets of the subgroup $\{I, P\}$ coincide, which is an indication of a normal subgroup).

The first relation (9.20) is obvious.

To prove the second relation (9.20), we make use of the following obvious properties of the reflection P :

$$PP = I, PaP = a, a \in O(3).$$

Multiplying the relation PaP on the left by P and using equation $PP = I$, we get the second relation (9.20).

Let us now prove the following assertion.

The subgroup $SO(3)$ of proper orthogonal transformations of the group $O(3)$ is isomorphic to the factor group $O(3)$ by the normal subgroup $\{I, P\}$.

Proof. The coset of the element $a \in SO(3)$ by the subgroup $\{I, P\}$ has the form $\{a, Pa\}$, Pa being an improper transformation (the product of the proper transformation a by the improper transformation P is an improper transformation).

If a' is an improper transformation, then the coset $\{a', Pa'\}$ can be reduced to the form $\{a, Pa\}$, where $a = Pa'$ is a proper transformation and $Pa = P(Pa') = (PP)a' = a'$.

Thus the factor group $O(3) / \{I, P\}$ consists of cosets of the form $\{a', Pa'\}$ where a is a proper transformation. The relation $a \leftrightarrow \{a, Pa\}$ is evidently an isomorphism between the groups $SO(3)$ and $O(3) / \{I, P\}$. We have proved the assertion.

9.2.6. 'The Lorentz group. We introduced in 8.4.1 the concept of a pseudo-Euclidean space $E_{(p,q)}^n$, i.e. a linear space in which a scalar product (x, y) equal to the non-singular symmetric bilinear form $A(x, y)$, the polar of the alternating quadratic form $A(x, x)$, is specified:

$$(x, y) = A(x, y). \quad (9-21)$$

We pointed out in 8.4.2 that in the so-called Galilean system of coordinates the square of the interval

$$s^2(x) = (x, x) \quad (9.22)$$

(it is an accepted term for the square of the length of the vector x with the coordinates $(x_1, x_2, \dots, x_n)$) had the form

$$s^2(x) = \sum_{i=1}^p x_i^2 - \sum_{i=p+1}^n x_i^2. \quad (9.23)$$

We must introduce the concept of the **Lorentz transformation** of the pseudo-Euclidean space $E_{(p,q)}^n$.

Definition. A linear transformation P of the pseudo-Euclidean space $E_{(p,q)}^n$ is called a **Lorentz transformation** if for any x and y from $E_{(p,q)}^n$ there holds a relation

$$(Px, Py) = (x, y), \quad (9.24)$$

where (x, y) is a scalar product defined by relation (9.21).

Equation (9.24) is a **condition for a Lorentz transformation**.

Note that under the Lorentz transformation the square of the interval $s^2(x)$, defined by relations (9.22) (or (9.23)) is retained.

As in 9.2.4, we can prove that the determinant $\det P$ of a Lorentz transformation is nonzero and, therefore, for every Lorentz transformation there is an inverse transformation P^{-1} .

In addition, by the meaning of the definition of a Lorentz transformation the product of such transformations is a Lorentz transformation. Thus, the following assertion holds true.

*The set of all Lorentz transformations of the pseudo-Euclidean space $E_{(p, q)}^n$ with an ordinary operation of multiplication of linear transformations (linear operators) forms a group which is called a **general Lorentz group** of the pseudo-Euclidean space $E_{(p, q)}^n$ and is designated as $L(n, p, q)$.*

We isolate a special class of pseudo-Euclidean spaces $E_{(1, n-1)}^n$ (it includes a space $E_{(1, 3)}^n$ which is interesting from the point of view of physics).

The Lorentz group for the spaces $E_{(1, n-1)}^n$ is designated as $L(n)$.

In 8.4.1 (formula (8.69)) we introduced the concept of the length $\sigma(x)$ of the vector x , which can be calculated by the formula

$$\sigma(x) = (\operatorname{sgn} s^2(x)) \sqrt{|s^2(x)|}.$$

By means of this formula all nonzero vectors of a pseudo-Euclidean space are divided into **time-like** ($\sigma(x) > 0$), **space-like** ($\sigma(x) < 0$) and isotropic ($\sigma(x) = 0$). It was proved that the set of the termini of the time-like, (space-like, isotropic) vectors whose origins coincide with an arbitrary fixed point, forms a cone T (the cone of space-like vectors is denoted by S). By agreement, the cone T is

divided into two connected components T^+ and T^- (the cone of the future and the cone of the past) so that each component, together with the vector $\mathbf{x}$, contains any vector $\lambda\mathbf{x}$, where $\lambda > 0$.

The division of vectors in a pseudo-Euclidean space just described makes it possible to separate some subgroups from the Lorentz group $L(n)$.

Namely, the subgroup $L(n)$ whose transformations carry any time-like vector again into a time-like vector is called a **complete Lorentz group**. It is designated as $L_{\uparrow}(n)$.

Another subgroup of the group $L(n)$ can be distinguished. It contains transformations the determinant of whose matrix is positive. This subgroup is designated as $L_+(n)$ and is called a **proper Lorentz group**.

The proper Lorentz transformations which belong to the subgroup $L_{\uparrow}(n)$ also form a subgroup. It is often called a **Lorentz group** and designated as $L_{\uparrow}(n)$.

To complete this subsection, we want to point out that *as distinct from orthogonal groups the Lorentz groups are noncompact*.

We shall prove, for example, the noncompactness of the group $L_{\uparrow}(2)$.

We completely described this group in 8.4.3. Recall that if a system of coordinates (x, y) has been introduced in $E_{(1,1)}^2$ such that the square of the interval is specified by the formula

$$s^2 = x^2 - y^2, \quad (9.25)$$

then the Lorentz transformations from the group $L_{\uparrow}(2)$ of the space $E_{(1,1)}^2$ are specified by the formulas

$$x' = \frac{x - \beta y}{\sqrt{1 - \beta^2}}, \quad y' = \frac{y - \beta x}{\sqrt{1 - \beta^2}}, \quad (9.26)$$

Let us consider a vector $\mathbf{x}$ with coordinates (0, 1) in the plane (x, y). By formula (9.26), this vector passes into a vector $\mathbf{x}_\beta$ with co-ordinates

$$\left(\frac{-\beta}{\sqrt{1-\beta^2}}, \frac{1}{\sqrt{1-\beta^2}} \right). \quad (9.27)$$

Let us consider now the sequence of the Lorentz transformations (9.26) defined by the values β_n from the relation

$$1 - \beta_n^2 = \frac{1}{n}, n = 1, 2, \dots \quad (9.28)$$

In accordance with (9.27) and (9.28), under the action of the indicated sequence of the Lorentz transformations, the vector $\mathbf{x}$ passes into the following sequence of the vectors $\{\mathbf{x}_n\}$ with coordinates

$$(-\sqrt{n-1}, \sqrt{n}). \quad (9.29)$$

Thus, for the values β from equation (9.28) we cannot isolate a convergent sequence from an infinite sequence of the Lorentz transformations in the group $L_{\uparrow}(2)$ defined by relations (9.26) (recall that the sequence of the elements A_n of a group is said to be convergent to the element A if the sequence $\{A_n \mathbf{x}\}$ converges to $A\mathbf{x}$ for any $\mathbf{x}$) since sequence (9.29) is unbounded.

The geometrical illustration of the noncompactness of the group $L_{\uparrow}(2)$ consists in the following.

In accordance with (9.25), a circle of radius unity in a pseudo-Euclidean plane is a hyperbola $x^2 - y^2 = 1$ which is a noncompact set. After a sequence of transformations from the group $L_{\uparrow}(2)$ considered above, the given point on this circle is transformed into an infinitely large sequence of points on the hyperbola indicated above, and we cannot isolate a convergent sequence from an infinitely large sequence of points.

9.2.7. Unitary groups. We shall consider here a complex linear space. By a complete analogy with 9.2.2, we can consider groups of linear transformations of that space. Since a complex number is defined by two real numbers (the real and the imaginary part), the complete linear group $GL(n)$ of transformation of an n -dimensional complex linear space is isomorphic to the complete linear group of transformations of the real $2n$ -dimensional space $GL(2n)$ (we can replace this symbol by $GL(2n, R)$ emphasizing that we speak of a group of transformations of a real space).

By analogy with a real Euclidean space, we regard, in a complete linear group of transformations of a complex Euclidean space, the so-called unitary groups $U(n)$ which are an analog of orthogonal groups (recall that in 5.7 unitary transformations (unitary operators) were defined as linear transformations retaining a scalar product; orthogonal transformations in a real case were defined in the same way).

As in a real case, we can isolate in the group $U(n)$ of unitary transformations a subgroup $SU(n)$ for which the determinants of unitary transformations are equal to unity.

9.3. Group Representations

In 9.2 we considered groups of linear transformations of a linear space. Thus, we investigated linear transformations from the point of view of their group properties. We have also considered geometrical and other properties of linear transformations.

We shall be interested here, in a definite sense, in a converse question: to what extent the properties of an abstractly defined group can be characterized by means of groups of linear transformations.

One of the means of solving this problem consists in the homomorphic (and, in particular, an isomorphic) mapping of an abstract group onto a subgroup (or the entire group) of linear transformations.

Thus, we arrive at the concept of **representation** of a given group by means of a subgroup of the group of linear transformations*.

The study of various representations of a given group makes it possible to elucidate significant properties of the group which are important for applications.

In many branches of physics (crystallography, the theory of relativity, quantum mechanics, etc) it is sometimes necessary to construct representations of various groups (finite and discrete subgroups of the group $GL(3)$, groups $O(3)$, $L(3)$, $U(3)$, $SU(3)$, etc). These constructions are far beyond the scope of the fundamental concepts of the group theory and cannot be discussed in this book.

We shall restrict our consideration to some concepts used in the theory of representations and to some examples.

9.3.1. Linear representations of groups. Terminology.

Definition. A *linear representation* of the group G in a finite-dimensional Euclidean space E^n is a mapping f which puts every element a of that group into correspondence with the linear transformation T_a of the space E^n such that the relation $T_{(a_1 a_2)} = T_{a_1} T_{a_2}$ holds for any a_1 and a_2 belonging to G .

Thus, the linear representation of the group G in the finite-dimensional Euclidean space E^n is a homomorphic mapping of that group onto a certain subset of the linear transformations of that space.

The following terminology is used: the space E^n is called the **space of representation**, the dimension n of that space is called the dimension of the representation, and the basis in the space E^n is called the **basis of the representation**.

Note that the homomorphic image $f(G)$ of the group G is also called the representation of that group in the space of representations.

In what follows, for the sake of brevity, n -dimensional linear representations of a group will be simply called the representations of that group.

The symbol $D(G)$ is used to designate the representation of the group G ; various representations of the given group are denoted by indices (say, $D^{(\mu)}(G)$). The symbol $D^{(\mu)}(g)$ is used for the linear transformation (linear operator) corresponding to the element $g \in G$ in the representation $D^{(\mu)}(G)$.

The **trivial representation of the group G** is a homomorphic mapping of G into a unit element of the group $GL(n)$.

If the mapping f of the group G onto the subgroup $GL(n)$ is isomorphic, the representation is called exact.

It is evident that not every group has an exact n -dimensional representation for a given n . For example, the group $O(11)$ cannot have an exact one-dimensional representation (this follows, in particular, from the fact that the group $O(1)$ is Abelian and the group $O(10)$ is not Abelian).

Note that the representation of a group resulting from the homomorphic mapping f of the group G into $GL(n)$ is isomorphic to the factor group $G / \text{kern } f$, where $\text{kern } f$ is the so-called congruence kernel of the homomorphism f , i.e. the set of the elements of G which is mapped into the unity of the group $GL(n)$ under the homomorphic mapping f .

9.3.2. Matrices of linear representations. Equivalent representations. Let us consider the representation $D^{(\mu)}(G)$ of the group G . In this representation, every element g from G is associated with the linear transformation $D^{(\mu)}(g)$. We designate the matrix of this linear transformation in the basis of the representation $D^{(\mu)}(G)$ as $D^{(\mu)i}(G)$ or $D_{ij}^{(\mu)}(g)$.

The matrix $D_{ij}^{(\mu)}(g)$ corresponding to the element g changes with the choice of the basis in the space of the representations. The question concerning the **equivalence of representations** of the group in the same space naturally arises.

Here is the definition of the equivalence of representations.

Definition. *The representations $D^{(\mu_1)}(G)$ and $D^{(\mu_2)}(G)$ of the group G in the same space E^n are said to be equivalent if there is a non-singular linear transformation C of the space E^n such that the relation $D^{(\mu_1)}(g) = C^{-1} D^{(\mu_2)}(g) C$ holds for every element $g \in G$.*

The concept of equivalence plays a significant role in the theory of representations, mainly in the enumeration and classification of representations.

The choice of a basis in the space of representations is also significant because in some basis the matrices corresponding to the elements of the group may have a standard, sufficiently simple form, which enables us to make important inferences concerning the representation under investigation. In 9.3.3 we shall give a certain classification of representations proceeding from a special form of matrices.

9.3.3. Reducible and irreducible representations. We shall discuss here the conditions under which the given representation $D(G)$, defined in the space E^n , induces a representation $\tilde{D}(G)$ in the subspace E' of that space.

This problem is closely connected with that of describing a given representation by means of more simple representations which have a smaller dimension than the given representation.

The concept of an invariant subspace of a linear transformation (linear operator) is closely connected with the solution of this problem.

Recall that the subspace E' is an invariant subspace of the linear operator A if for every element x from E' the element Ax belongs to E' (see 5.3). In other words, the subspace E' is invariant if the action of the operator A on the elements of that subspace does not take them outside that subspace. Note that the space E^n itself and the zero element of the space are invariant subspaces of any linear operator.

We can introduce the concept of an invariant subspace for the representation $D(G)$. *Namely, the subspace E' is said to be invariant for the representation $D(G)$ if it is invariant for every operator from $D(G)$.*

It is evident that a certain representation $\tilde{D}(G)$ is induced on the invariant subspace of the representation $D(G)$. It should be pointed out that the representation $\tilde{D}(G)$ does not reduce to the representation $D(G)$ if the invariant subspace E' does not coincide with E^n .

Let us discuss the concept of a **reducible** representation in more detail.

Suppose, for instance, that all the matrices of a certain three-dimensional representation $D(G)$ have the form

$$\left(\begin{array}{c|c} A_1 & A_2 \\ \hline 0 & A_3 \end{array} \right) = \left(\begin{array}{cc|c} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ \hline 0 & 0 & a_{33} \end{array} \right), \quad (9.30)$$

where A_1, A_2, A_3 and O denote the matrices $\begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$, $\begin{pmatrix} a_{13} \\ a_{23} \end{pmatrix}$, (a_{33}) and $(0, 0)$ respectively. It is easy to verify that the product of the matrices of form (9.30) obeys the law

$$\left(\begin{array}{c|c} A'_1 & A'_2 \\ \hline 0 & A'_3 \end{array} \right) \left(\begin{array}{c|c} A''_1 & A''_2 \\ \hline 0 & A''_3 \end{array} \right) = \left(\begin{array}{c|c} A'_1 A''_1 & A'_2 A''_2 \\ \hline 0 & A'_3 A''_3 \end{array} \right),$$

i.e. the product of the matrices of form (9.30) is a matrix of the form (9.30). Moreover, when matrices of this form are multiplied, the matrices A'_1 and A''_1 and the matrices A'_3 and A''_3 are multiplied between themselves.

Thus, we see that the matrix $A = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$ forms a two-dimensional representation of the group being considered, and the matrix $A = (a_{33})$ forms a one-dimensional representation of the same group.

In these cases we say that the representation $D(G)$ is **reducible**.

If all the matrices (we speak of square matrices of order n) of the operators of the representation have the form

$$\left(\begin{array}{c|c} A_1 & 0 \\ \hline 0 & A_2 \end{array} \right), \quad (9.31)$$

where A_1 and A_2 are square matrices of different orders in a general case, then it is clear that the matrices A_1 and A_2 form representations the sum of whose dimensions is equal to n .

In this case the representation is **fully reducible**. Note that the operators whose matrices have the form (9.31) are actually reduced to two operators acting independently in two invariant subspaces.

It should also be pointed out that the representation which is induced on an invariant subspace by the given representation $D(G)$ is called a part of the representation $D(G)$.

In conclusion, we shall formulate the concept of an **irreducible** representation.

*The representation $D(G)$ of the group G is said to be **irreducible** if it has only two invariant subspaces, E^n and O .*

Otherwise, the representation is **reducible**.

The part played by irreducible representations consists in the fact that any representation can be expressed in terms of irreducible representations.

9.3.4. Characters. In the theory of representations of groups and, in particular, in the theory of representations of finite groups a useful role is played by the invariants of linear transformations which form a representation. The importance of the invariants is also clear from the fact that they do not depend on the choice of the basis of a representation and, therefore, they characterize the representation in a certain sense. .

Assume that $D(G)$ is an n -dimensional representation of the group G , and $D_j^i(g)$ is the matrix of the operator corresponding to the element g from G .

*The **character** of the element $g \in G$ in the representation $D(G)$ is the number*

$$\chi(g) = D_i^i(g) = D_1^1(g) + D_2^2(g) + \dots + D_n^n(g).$$

Thus, the character of the element g is the trace of the matrix of the operator $D(g)$.

Since the trace of the matrix of a linear operator is an invariant (see 5.2.3), the character of any element does not depend on the basis of the representation and is, therefore, invariant.

Thus, every element $g \in G$ of the representation $D(G)$ is associated with a number which is the character of that element.

Since different elements may have the same characters, we must find which elements of the group are associated with the same characters. To solve this problem, we introduce the concept of conjugate elements and the classes of conjugate elements in the given group G .

*The element $b \in G$ is said to be the **conjugate** element of $a \in G$ if there is an element $u \in G$ such that*

$$uau^{-1} = b. \quad (9.32)$$

Note the following properties of conjugate elements.

(1) *Every element a is the conjugate of itself.* Indeed, if e is the unity of a group, then, evidently, the relation $ea e^{-1} = a$ holds true, which means that a is the conjugate element of a .

(2) *If the element b is the conjugate of a , then the element a is the conjugate of b .* This property follows immediately from (9.32). Indeed, multiplying both sides of (9.32) u^{-1} on the left and by u on the right, we find that $u^{-1}bu = a$. Noting that the inverse element of u^{-1} is u , we verify the validity of the property we have formulated.

(3) *If b is the conjugate element of a and c is the conjugate element of b , then c is the conjugate of a .*

Indeed, since $c = vbv^{-1}$, and $b = uau^{-1}$, then, evidently,

$$c = vub u^{-1} v^{-1}. \quad (9.33)$$

Since the inverse element of vu is $u^{-1}v^{-1}$, it follows from (9.33) that the element c is the conjugate of a .

Let us place in the same class all the elements of the group which are the conjugates of the given element a . Thus, according to property (3) every element of the class is the conjugate of any element of that class. Two classes of this kind evidently either coincide or have no elements in common.

Let us return now to group representations.

Assume that a and b are conjugate elements, i.e. relation (9.32) holds true:

$$b = uau^{-1}. \quad (9.32)$$

Let us direct our attention to the operators $D(a)$, $D(b)$, $D(u)$ and $D(u^{-1})$. According to the definition of a group representation, the operator $D(u^{-1})$ is an inverse of the operator $D(u)$, i.e. $D(u^{-1}) = (D(u))^{-1}$.

Using again the definition of a representation, we get, in accordance with (9.32), a relation $D(b) = D(u) D(a) (D(u))^{-1}$.

Let us pass now to the matrices of the operators appearing in the last relation. We see that we can regard the matrix of the operator $D(b)$ as that of the operator $D(a)$ when we pass to a new basis with the transition matrix $D(u)$. (see 5.2.2). 'Since under such transformations the trace of the matrix is invariant and by definition is equal to the character of the element, we can infer that $\chi(a) = \chi(b)$.

Thus, the characters of all elements belonging to the same class of conjugate elements are equal to each other.

It is also evident that *the characters of the elements for equivalent representations coincide.*

The concept of a character is usually used in the theory of representations as follows.

Assume that the given group G can be divided into a finite number of various classes of the conjugate elements $K_1, K_2, \dots, K_v$.

Then every element of the class K_i in the given representation $D(G)$ (and in any equivalent representation) is associated with the same character χ_i . Therefore, the representation $D(G)$ can be described by means of a collection of characters $\chi_1, \chi_2, \dots, \chi_v$ which can be regarded as the coordinates of a vector in a Euclidean space of dimension v . Thus different vectors correspond to different representations. The indicated geometrical approach makes it possible to solve important problems of the group representation theory in many cases.

9.3.5. Examples of group representations.

Example 1. Assume that G is a symmetry group of a three-dimensional space consisting of two elements: of an identity transformation I (the unity of the group) and the reflection P about the origin. Thus $G = \{I, P\}$.

The multiplication of the elements of the group is defined by the following table:

	I	P
I	I	P
P	P	I

(9.34)

(1) *A one-dimensional representation of the group G .*

We choose a basis e_1 in the space E^1 and consider the matrix $A^{(1)}$ of the linear non-singular transformation $A^{(1)}$ in that space: $A^{(1)} = 1$. transformation $A^{(1)}$ evidently forms a subgroup in the group $GL(1)$ of the linear transformations of the space E^1 , the multiplication in this subgroup being defined by the table

	$A^{(1)}$
$A^{(1)}$	$A^{(1)}$

We evidently get a one-dimensional representation $D^{(1)}(G)$ of the group G by means of the relations $D^{(1)}(I) = A^{(1)}$, $D^{(1)}(P) = A^{(1)}$ (these relations define the homomorphism of the group G into $GL(1)$ and, consequently, its representation).

(2) *A two-dimensional representation of the group G .*

We choose some basis e_1, e_2 in E^2 and consider in that basis the matrices $A^{(2)}$ and $B^{(2)}$ of linear non-singular transformations $A^{(1)}$ and $B^{(2)}$:

$$A^2 = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}, \quad B^2 = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$$

(since $\det A^{(2)} = 1$ and $\det B^{(2)} = -1$, are nonsingular transformations).

The transformations $A^{(2)}$ and $B^{(2)}$ form a subgroup in the group $GL(2)$. We make sure by a direct verification (by multiplying the matrices $A^{(2)}$ and $B^{(2)}$) that the multiplication of the operators $A^{(2)}$ and $B^{(2)}$ is defined by the table

	$A^{(2)}$	$B^{(2)}$
$A^{(2)}$	$A^{(2)}$	$B^{(2)}$
$B^{(2)}$	$B^{(2)}$	$A^{(2)}$

(9.35)

We obtain a two-dimensional representation $D^{(2)}(G)$ of the group G by the relations

$$D^{(2)}(I) = A^{(2)}, \quad D^{(2)}(P) = B^{(2)}. \quad (9.36)$$

Indeed, comparing tables (9.34 and (9.35), we see that (9.36) defines the isomorphism of the group G into the subgroup $\{A^{(2)}, B^{(2)}\}$ of the group $GL(2)$ and, consequently, the representation of this group.

(3) *A three-dimensional representation of the group G .*

Let us consider in E^3 a linear representation $A^{(3)}$ defined by the matrix

$$A^{(3)} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

This transformation forms a subgroup of the group $GL(3)$ with the multiplication law $A^{(3)} A^{(3)} = A^{(3)}$.

As in the case of a one-dimensional representation, we get a three-dimensional representation $D^{(3)}(G)$ by means of the relations

$$D^{(3)}(I) = A^{(3)}, \quad D^{(3)}(P) = A^{(3)}.$$

(4) *A four-dimensional representation of the group G .*

Let us consider in E^4 the linear transformations $A^{(4)}$ and $B^{(4)}$ defined by the matrices

$$A^{(4)} = \left(\begin{array}{cc|cc} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ \hline 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{array} \right), \quad B^{(4)} = \left(\begin{array}{cc|cc} 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ \hline 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{array} \right).$$

The transformations $A^{(4)}$ and $B^{(4)}$ form a subgroup in the group $GL(4)$, with the multiplication law defined by a table similar to (9.35) (with the index 2 replaced by 4). We evidently obtain a four-dimensional representation $D^{(4)}(G)$ of the group G by the relations

$$D^{(4)}(I) = A^{(4)}, \quad D^{(4)}(P) = B^{(4)}.$$

Remark. It is easy to see that the matrices $A^{(4)}$ and $B^{(4)}$ can be written in the form

$$A^{(4)} = \begin{pmatrix} A^{(2)} & 0 \\ 0 & A^{(2)} \end{pmatrix}, \quad B^{(4)} = \begin{pmatrix} B^{(2)} & 0 \\ 0 & B^{(2)} \end{pmatrix}.$$

The representation $D^{(4)}(G)$ can, therefore, be conventionally written in the form $D^{(4)}(G) = D^{(2)}(G) + D^{(2)}(G) = 2D^{(2)}(G)$. Quite similarly we can conventionally write $D^{(3)}(G)$ in the form $D^{(3)}(G) = 3D^{(1)}(G)$.

Using this remark, the reader will easily construct a representation of the group G of any finite dimensionality.

Example 2. We have proved in 9.2.5 that the symmetry group $G = \{I, P\}$ of a three-dimensional space just considered was a normal subgroup of the group $O(3)$ (a group of orthogonal transformations of the space E^3). We proved there that the subgroup $SO(3)$ of the proper orthogonal transformations of the group $O(3)$ was isomorphic to the factor group of the group $O(3)$ by the normal subgroup $\{I, P\}$.

Since the group is homomorphically mapped onto each of its factor groups, it follows that $O(3)$ can be homomorphically mapped onto the group $SO(3)$. As we saw in 9.3.5, the indicated homomorphism was carried out as follows.

If a is a proper transformation from $O(3)$, it is associated with the same transformation from $SO(3)$.

If a' is an improper transformation, then it is associated with a proper transformation Pa' .

We thus obtain a three-dimensional representation $DO(3)$ of the group of orthogonal transformations by means of the group $SO(3)$ of proper orthogonal transformations.