

Contents

1	Vector spaces	3
1.1	Definitions	3
1.2	Bases	5
1.3	Row and column vectors	9
1.4	Change of basis	11
1.5	Subspaces and direct sums	13
2	Matrices and determinants	15
2.1	Matrix algebra	15
2.2	Row and column operations	16
2.3	Rank	20
2.4	Determinants	22
2.5	Calculating determinants	25
2.6	The Cayley–Hamilton Theorem	29
3	Linear maps between vector spaces	33
3.1	Definition and basic properties	33
3.2	Representation by matrices	35
3.3	Change of basis	37
3.4	Canonical form revisited	39
4	Linear maps on a vector space	41
4.1	Projections and direct sums	41
4.2	Linear maps and matrices	43
4.3	Eigenvalues and eigenvectors	44
4.4	Diagonalisability	45
4.5	Characteristic and minimal polynomials	48
4.6	Jordan form	51
4.7	Trace	52

- (a) The vectors $v_1, v_2, \dots, v_n$ are *linearly independent* if, whenever we have scalars $c_1, c_2, \dots, c_n$ satisfying

$$c_1 v_1 + c_2 v_2 + \dots + c_n v_n = 0,$$

then necessarily $c_1 = c_2 = \dots = 0$.

- (b) The vectors $v_1, v_2, \dots, v_n$ are *spanning* if, for every vector $v \in V$, we can find scalars $c_1, c_2, \dots, c_n \in \mathbb{K}$ such that

$$v = c_1 v_1 + c_2 v_2 + \dots + c_n v_n.$$

In this case, we write $V = \langle v_1, v_2, \dots, v_n \rangle$.

- (c) The vectors $v_1, v_2, \dots, v_n$ form a *basis* for V if they are linearly independent and spanning.

Remark Linear independence is a property of a *list* of vectors. A list containing the zero vector is never linearly independent. Also, a list in which the same vector occurs more than once is never linearly independent.

I will say “Let $B = (v_1, \dots, v_n)$ be a basis for V ” to mean that the list of vectors $v_1, \dots, v_n$ is a basis, and to refer to this list as B .

Definition 1.4 Let V be a vector space over the field $\mathbb{K}$. We say that V is *finite-dimensional* if we can find vectors $v_1, v_2, \dots, v_n \in V$ which form a basis for V .

Remark In these notes we are only concerned with finite-dimensional vector spaces. If you study Functional Analysis, Quantum Mechanics, or various other subjects, you will meet vector spaces which are not finite dimensional.

Proposition 1.1 The following three conditions are equivalent for the vectors $v_1, \dots, v_n$ of the vector space V over $\mathbb{K}$:

- (a) $v_1, \dots, v_n$ is a basis;
- (b) $v_1, \dots, v_n$ is a maximal linearly independent set (that is, if we add any vector to the list, then the result is no longer linearly independent);
- (c) $v_1, \dots, v_n$ is a minimal spanning set (that is, if we remove any vector from the list, then the result is no longer spanning).

The next theorem helps us to understand the properties of linear independence.

Theorem 1.2 (The Exchange Lemma) *Let V be a vector space over $\mathbb{K}$. Suppose that the vectors $v_1, \dots, v_n$ are linearly independent, and that the vectors $w_1, \dots, w_m$ are linearly independent, where $m > n$. Then we can find a number i with $1 \leq i \leq m$ such that the vectors $v_1, \dots, v_n, w_i$ are linearly independent.*

In order to prove this, we need a lemma about systems of equations.

Lemma 1.3 *Given a system (*)*

$$\begin{aligned} a_{11}x_1 + a_{12}x_2 + \dots + a_{1m}x_m &= 0, \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2m}x_m &= 0, \\ &\dots \\ a_{n1}x_1 + a_{n2}x_2 + \dots + a_{nm}x_m &= 0 \end{aligned}$$

of homogeneous linear equations, where the number n of equations is strictly less than the number m of variables, there exists a non-zero solution $(x_1, \dots, x_m)$ (that is, $x_1, \dots, x_m$ are not all zero).

Proof This is proved by induction on the number of variables. If the coefficients $a_{11}, a_{21}, \dots, a_{n1}$ of x_1 are all zero, then putting $x_1 = 1$ and the other variables zero gives a solution. If one of these coefficients is non-zero, then we can use the corresponding equation to express x_1 in terms of the other variables, obtaining $n - 1$ equations in $m - 1$ variables. By hypothesis, $n - 1 < m - 1$. So by the induction hypothesis, these new equations have a non-zero solution. Computing the value of x_1 gives a solution to the original equations.

Now we turn to the proof of the Exchange Lemma. Let us argue for a contradiction, by assuming that the result is false: that is, assume that none of the vectors w_i can be added to the list $(v_1, \dots, v_n)$ to produce a larger linearly independent list. This means that, for all j , the list $(v_1, \dots, v_n, w_j)$ is linearly dependent. So there are coefficients $c_1, \dots, c_n, d$, not all zero, such that

$$c_1v_1 + \dots + c_nv_n + dw_i = 0.$$

We cannot have $d = 0$; for this would mean that we had a linear combination of $v_1, \dots, v_n$ equal to zero, contrary to the hypothesis that these vectors are linearly independent. So we can divide the equation through by d , and take w_i to the other side, to obtain (changing notation slightly)

$$w_i = a_{1i}v_1 + a_{2i}v_2 + \dots + a_{ni}v_n = \sum_{j=1}^n a_{ji}v_j.$$

Remark We allow the possibility that a vector space has dimension zero. Such a vector space contains just one vector, the zero vector 0 ; a basis for this vector space consists of the empty set.

Now let V be an n -dimensional vector space over $\mathbb{K}$. This means that there is a basis $v_1, v_2, \dots, v_n$ for V . Since this list of vectors is spanning, every vector $v \in V$ can be expressed as

$$v = c_1 v_1 + c_2 v_2 + \cdots + c_n v_n$$

for some scalars $c_1, c_2, \dots, c_n \in \mathbb{K}$. The scalars $c_1, \dots, c_n$ are the *coordinates* of v (with respect to the given basis), and the *coordinate representation* of v is the n -tuple

$$(c_1, c_2, \dots, c_n) \in \mathbb{K}^n.$$

Now *the coordinate representation is unique*. For suppose that we also had

$$v = c'_1 v_1 + c'_2 v_2 + \cdots + c'_n v_n$$

for scalars $c'_1, c'_2, \dots, c'_n$. Subtracting these two expressions, we obtain

$$0 = (c_1 - c'_1)v_1 + (c_2 - c'_2)v_2 + \cdots + (c_n - c'_n)v_n$$

Now the vectors $v_1, v_2, \dots, v_n$ are linearly independent; so this equation implies that $c_1 - c'_1 = 0$, $c_2 - c'_2 = 0$, $\dots$, $c_n - c'_n = 0$; that is,

$$c_1 = c'_1, \quad c_2 = c'_2, \quad \dots, \quad c_n = c'_n.$$

Now it is easy to check that, when we add two vectors in V , we add their coordinate representations in $\mathbb{K}^n$ (using coordinatewise addition); and when we multiply a vector $v \in V$ by a scalar c , we multiply its coordinate representation by c . In other words, addition and scalar multiplication in V translate to the same operations on their coordinate representations. This is why we only need to consider vector spaces of the form $\mathbb{K}^n$, as in Example 1.2.

Here is how the result would be stated in the language of abstract algebra:

Theorem 1.5 *Any n -dimensional vector space over a field $\mathbb{K}$ is isomorphic to the vector space $\mathbb{K}^n$.*

1.3 Row and column vectors

The elements of the vector space $\mathbb{K}^n$ are all the n -tuples of scalars from the field $\mathbb{K}$. There are two different ways that we can represent an n -tuple: as a row, or as

In matrix form, this says

$$\begin{bmatrix} c_1 \\ \vdots \\ c_n \end{bmatrix} = P \begin{bmatrix} d_1 \\ \vdots \\ d_n \end{bmatrix},$$

or in other words

$$[v]_B = P[v]_{B'},$$

as required.

In this course, we will see four ways in which matrices arise in linear algebra. Here is the first occurrence: **matrices arise as transition matrices between bases of a vector space.**

The next corollary summarises how transition matrices behave. Here I denotes the *identity matrix*, the matrix having 1s on the main diagonal and 0s everywhere else. Given a matrix P , we denote by P^{-1} the *inverse* of P , the matrix Q satisfying $PQ = QP = I$. Not every matrix has an inverse: we say that P is *invertible* or *non-singular* if it has an inverse.

Corollary 1.7 Let B, B', B'' be bases of the vector space V .

(a) $P_{B,B} = I$.

(b) $P_{B',B} = (P_{B,B'})^{-1}$.

(c) $P_{B,B''} = P_{B,B'} P_{B',B''}$.

This follows from the preceding Proposition. For example, for (b) we have

$$[v]_B = P_{B,B'} [v]_{B'}, \quad [v]_{B'} = P_{B',B} [v]_B,$$

so

$$[v]_B = P_{B,B'} P_{B',B} [v]_B.$$

By the uniqueness of the coordinate representation, we have $P_{B,B'} P_{B',B} = I$.

Corollary 1.8 The transition matrix between any two bases of a vector space is invertible.

This follows immediately from (b) of the preceding Corollary.

Remark We see that, to express the coordinate representation w.r.t. the new basis in terms of that w.r.t. the old one, we need the inverse of the transition matrix:

$$[v]_{B'} = P_{B,B'}^{-1} [v]_B.$$

Theorem 2.1 Let A be an $m \times n$ matrix over the field $\mathbb{K}$. Then it is possible to change A into B by elementary row and column operations, where B is a matrix of the same size satisfying $B_{ii} = 1$ for $0 \leq i \leq r$, for $r \leq \min\{m, n\}$, and all other entries of B are zero.

If A can be reduced to two matrices B and B' both of the above form, where the numbers of non-zero elements are r and r' respectively, by different sequences of elementary operations, then $r = r'$, and so $B = B'$.

Definition 2.4 The number r in the above theorem is called the *rank* of A ; while a matrix of the form described for B is said to be in the *canonical form for equivalence*. We can write the canonical form matrix in “block form” as

$$B = \begin{bmatrix} I_r & O \\ O & O \end{bmatrix},$$

where I_r is an $r \times r$ identity matrix and O denotes a zero matrix of the appropriate size (that is, $r \times (n-r)$, $(m-r) \times r$, and $(m-r) \times (n-r)$ respectively for the three O s). Note that some or all of these O s may be missing: for example, if $r = m$, we just have $[I_m \ O]$.

Proof We outline the proof that the reduction is possible. To prove that we always get the same value of r , we need a different argument.

The proof is by induction on the size of the matrix A : in other words, we assume as inductive hypothesis that any smaller matrix can be reduced as in the theorem. For the matrix A we give, we proceed in steps as follows:

- If $A = O$ (the all-zero matrix), then the conclusion of the theorem holds, with $r = 0$; no reduction is required. So assume that $A \neq O$.
- If $A_{11} \neq 0$, then skip this step. If $A_{11} = 0$, then there is a non-zero element A_{ij} somewhere in A ; by swapping the first and i th rows, and the first and j th columns, if necessary (Type 3 operations), we can bring this entry into the $(1, 1)$ position.
- Now we can assume that $A_{11} \neq 0$. Multiplying the first row by A_{11}^{-1} , (row operation Type 2), we obtain a matrix with $A_{11} = 1$.
- Now by row and column operations of Type 1, we can assume that all the other elements in the first row and column are zero. For if $A_{1j} \neq 0$, then subtracting A_{1j} times the first column from the j th gives a matrix with $A_{1j} = 0$. Repeat this until all non-zero elements have been removed.

corresponds to the elementary column operation of adding twice the first column to the second, or to the elementary row operation of adding twice the second row to the first. For the other types, the matrices for row operations and column operations are identical.

Lemma 2.2 *The effect of an elementary row operation on a matrix is the same as that of multiplying on the left by the corresponding elementary matrix. Similarly, the effect of an elementary column operation is the same as that of multiplying on the right by the corresponding elementary matrix.*

The proof of this lemma is somewhat tedious calculation.

Example 2.3 We continue our previous example. In order, here is the list of elementary matrices corresponding to the operations we applied to A . (Here 2×2 matrices are row operations while 3×3 matrices are column operations).

$$\begin{bmatrix} 1 & -2 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}, \begin{bmatrix} 1 & 0 & -3 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}, \begin{bmatrix} 1 & 0 \\ -4 & 1 \end{bmatrix}, \begin{bmatrix} 1 & 0 \\ 0 & -1/3 \end{bmatrix}, \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & -2 \\ 0 & 0 & 1 \end{bmatrix}.$$

So the whole process can be written as a matrix equation:

$$\begin{bmatrix} 1 & 0 \\ 0 & -1/3 \end{bmatrix} \begin{bmatrix} 1 & 0 \\ -4 & 1 \end{bmatrix} A \begin{bmatrix} 1 & -2 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 1 & 0 & -3 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & -2 \\ 0 & 0 & 1 \end{bmatrix} = B,$$

or more simply

$$\begin{bmatrix} 1 & 0 \\ 4/3 & -1/3 \end{bmatrix} A \begin{bmatrix} 1 & 2 & 1 \\ 0 & 1 & -2 \\ 0 & 0 & 1 \end{bmatrix} = B,$$

where, as before,

$$A = \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{bmatrix}, \quad B = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \end{bmatrix}.$$

An important observation about the elementary operations is that each of them can have its effect undone by another elementary operation of the same kind, and hence every elementary matrix is invertible, with its inverse being another elementary matrix of the same kind. For example, the effect of adding twice the first row to the second is undone by adding -2 times the first row to the second, so that

$$\begin{bmatrix} 1 & 2 \\ 0 & 1 \end{bmatrix}^{-1} = \begin{bmatrix} 1 & -2 \\ 0 & 1 \end{bmatrix}.$$

Since the product of invertible matrices is invertible, we can state the above theorem in a more concise form. First, one more definition:

Definition 2.5 The $m \times n$ matrices A and B are said to be *equivalent* if $B = PAQ$, where P and Q are invertible matrices of sizes $m \times m$ and $n \times n$ respectively.

Theorem 2.3 Given any $m \times n$ matrix A , there exist invertible matrices P and Q of sizes $m \times m$ and $n \times n$ respectively, such that PAQ is in the canonical form for equivalence.

Remark The relation “equivalence” defined above is an equivalence relation on the set of all $m \times n$ matrices; that is, it is reflexive, symmetric and transitive.

When mathematicians talk about a “canonical form” for an equivalence relation, they mean a set of objects which are representatives of the equivalence classes: that is, every object is equivalent to a unique object in the canonical form. We have shown this for the relation of equivalence defined earlier, except for the uniqueness of the canonical form. This is our job for the next section.

2.3 Rank

We have the unfinished business of showing that the rank of a matrix is well defined; that is, no matter how we do the row and column reduction, we end up with the same canonical form. We do this by defining two further kinds of rank, and proving that all three are the same.

Definition 2.6 Let A be an $m \times n$ matrix over a field $\mathbb{K}$. We say that the *column rank* of A is the maximum number of linearly independent columns of A , while the *row rank* of A is the maximum number of linearly independent rows of A . (We regard columns or rows as vectors in $\mathbb{K}^m$ and $\mathbb{K}^n$ respectively.)

Now we need a sequence of four lemmas.

Lemma 2.4 (a) Elementary column operations don't change the column rank of a matrix.

(b) Elementary row operations don't change the column rank of a matrix.

(c) Elementary column operations don't change the row rank of a matrix.

(d) Elementary row operations don't change the row rank of a matrix.

Proof (a) This is clear for Type 3 operations, which just rearrange the vectors. For Types 1 and 2, we have to show that such an operation cannot take a linearly independent set to a linearly dependent set; the *vice versa* statement holds because the inverse of an elementary operation is another operation of the same kind.

So suppose that $v_1, \dots, v_n$ are linearly independent. Consider a Type 1 operation involving adding c times the j th column to the i th; the new columns are $v'_1, \dots, v'_n$, where $v'_k = v_k$ for $k \neq i$, while $v'_i = v_i + cv_j$. Suppose that the new vectors are linearly dependent. Then there are scalars $a_1, \dots, a_n$, not all zero, such that

$$\begin{aligned} 0 &= a_1 v'_1 + \dots + a_n v'_n \\ &= a_1 v_1 + \dots + a_i(v_i + cv_j) + \dots + a_j v_j + \dots + a_n v_n \\ &= a_1 v_1 + \dots + a_i v_i + \dots + (a_j + ca_i)v_j + \dots + a_n v_n. \end{aligned}$$

Since $v_1, \dots, v_n$ are linearly independent, we conclude that

$$a_1 = 0, \dots, a_i = 0, \dots, a_j + ca_i = 0, \dots, a_n = 0,$$

from which we see that all the a_k are zero, contrary to assumption. So the new columns are linearly independent.

The argument for Type 2 operations is similar but easier.

(b) It is easily checked that, if an elementary row operation is applied, then the new vectors satisfy exactly the same linear relations as the old ones (that is, the same linear combinations are zero). So the linearly independent sets of vectors don't change at all.

(c) Same as (b), but applied to rows.

(d) Same as (a), but applied to rows.

Theorem 2.5 For any matrix A , the row rank, the column rank, and the rank are all equal. In particular, the rank is independent of the row and column operations used to compute it.

Proof Suppose that we reduce A to canonical form B by elementary operations, where B has rank r . These elementary operations don't change the row or column rank, by our lemma; so the row ranks of A and B are equal, and their column ranks are equal. But it is trivial to see that, if

$$B = \begin{bmatrix} I_r & O \\ O & O \end{bmatrix},$$

then the row and column ranks of B are both equal to r . So the theorem is proved.

We can get an extra piece of information from our deliberations. Let A be an invertible $n \times n$ matrix. Then the canonical form of A is just I : its rank is equal to n . This means that there are matrices P and Q , each a product of elementary matrices, such that

$$PAQ = I_n.$$

In symbols,

$$(c_i, c_j) \mapsto (c_i, c_j + c_i) \mapsto (-c_i, c_j + c_i) \mapsto (c_j, c_j + c_i) \mapsto (c_j, c_i).$$

The first, third and fourth steps don't change the value of D , while the second multiplies it by -1 .

Now we take the matrix A and apply elementary column operations to it, keeping track of the factors by which D gets multiplied according to rules (a)–(c). The overall effect is to multiply $D(A)$ by a certain non-zero scalar c , depending on the operations.

- If A is invertible, then we can reduce A to the identity, so that $cD(A) = D(I) = 1$, whence $D(A) = c^{-1}$.
- If A is not invertible, then its column rank is less than n . So the columns of A are linearly dependent, and one column can be written as a linear combination of the others. Applying axiom (D1), we see that $D(A)$ is a linear combination of values $D(A')$, where A' are matrices with two equal columns; so $D(A') = 0$ for all such A' , whence $D(A) = 0$.

This proves that the determinant function, if it exists, is unique. We show its existence in the next section, by giving a couple of formulae for it.

Given the uniqueness of the determinant function, we now denote it by $\det(A)$ instead of $D(A)$. The proof of the theorem shows an important corollary:

Corollary 2.9 *A square matrix is invertible if and only if $\det(A) \neq 0$.*

Proof See the case division at the end of the proof of the theorem.

One of the most important properties of the determinant is the following.

Theorem 2.10 *If A and B are $n \times n$ matrices over $\mathbb{K}$, then $\det(AB) = \det(A) \det(B)$.*

Proof Suppose first that B is not invertible. Then $\det(B) = 0$. Also, AB is not invertible. (For, suppose that $(AB)^{-1} = X$, so that $XAB = I$. Then XA is the inverse of B .) So $\det(AB) = 0$, and the theorem is true.

In the other case, B is invertible, so we can apply a sequence of elementary column operations to B to get to the identity. The effect of these operations is to multiply the determinant by a non-zero factor c (depending on the operations), so that $c \det(B) = 1$, or $c = (\det(B))^{-1}$. Now these operations are represented by elementary matrices; so we see that $BQ = I$, where Q is a product of elementary matrices.

Spanning: Take any vector in $\text{Im}(\alpha)$, say w . Then $w = \alpha(v)$ for some $v \in V$. Write v in terms of the basis for V :

$$v = a_1 u_1 + \cdots + a_q u_q + c_1 v_1 + \cdots + c_s v_s$$

for some $a_1, \dots, a_q, c_1, \dots, c_s$. Applying α , we get

$$\begin{aligned} w &= \alpha(v) \\ &= a_1 \alpha(u_1) + \cdots + a_q \alpha(u_q) + c_1 \alpha(v_1) + \cdots + c_s \alpha(v_s) \\ &= c_1 w_1 + \cdots + c_s w_s, \end{aligned}$$

since $\alpha(u_i) = 0$ (as $u_i \in \text{Ker}(\alpha)$) and $\alpha(v_i) = w_i$. So the vectors $w_1, \dots, w_s$ span $\text{Im}(\alpha)$.

Thus, $\rho(\alpha) = \dim(\text{Im}(\alpha)) = s$. Since $\nu(\alpha) = q$ and $q + s = \dim(V)$, the theorem is proved.

3.2 Representation by matrices

We come now to the second role of matrices in linear algebra. They represent linear maps between vector spaces.

Let $\alpha : V \rightarrow W$ be a linear map, where $\dim(V) = m$ and $\dim(W) = n$. As we saw in the first section, we can take V and W in their coordinate representation: $V = \mathbb{K}^m$ and $W = \mathbb{K}^n$. The elements of these vector spaces being represented as column vectors. Let $e_1, \dots, e_m$ be the standard basis for V (so that e_i is the vector with i th coordinate 1 and all other coordinates zero), and $f_1, \dots, f_n$ the standard basis for W . Then for $i = 1, \dots, m$, the vector $\alpha(e_i)$ belongs to W , so we can write it as a linear combination of $f_1, \dots, f_n$.

Definition 3.4 The matrix representing the linear map $\alpha : V \rightarrow W$ relative to the bases $B = (e_1, \dots, e_m)$ for V and $C = (f_1, \dots, f_n)$ for W is the $n \times m$ matrix whose (i, j) entry is a_{ij} , where

$$\alpha(e_i) = \sum_{j=1}^n a_{ji} f_j$$

for $j = 1, \dots, n$.

In practice this means the following. Take $\alpha(e_i)$ and write it as a column vector $[a_{1i} \ a_{2i} \ \cdots \ a_{ni}]^T$. This vector is the i th column of the matrix representing α . So, for example, if $m = 3$, $n = 2$, and

$$\alpha(e_1) = f_1 + f_2, \quad \alpha(e_2) = 2f_1 + 5f_2, \quad \alpha(e_3) = 3f_1 - f_2,$$

Proposition 3.7 *Two matrices represent the same linear map with respect to different bases if and only if they are equivalent.*

This holds because

- transition matrices are always invertible (the inverse of $P_{B,B'}$ is the matrix $P_{B',B}$ for the transition in the other direction); and
- any invertible matrix can be regarded as a transition matrix: for, if the $n \times n$ matrix P is invertible, then its rank is n , so its columns are linearly independent, and form a basis B' for $\mathbb{K}^n$; and then $P = P_{B,B'}$, where B is the “standard basis”.

3.4 Canonical form revisited

Now we can give a simpler proof of Theorem 2.3 about canonical form for equivalence. First, we make the following observation.

Theorem 3.8 *Let $\alpha : V \rightarrow W$ be a linear map of rank $r = \rho(\alpha)$. Then there are bases for V and W such that the matrix representing α is, in block form,*

$$\begin{bmatrix} I_r & O \\ O & O \end{bmatrix}.$$

Proof As in the proof of Theorem 3.2, choose a basis $u_1, \dots, u_s$ for $\text{Ker}(\alpha)$, and extend to a basis $u_1, \dots, u_s, v_1, \dots, v_r$ for V . Then $\alpha(v_1), \dots, \alpha(v_r)$ is a basis for $\text{Im}(\alpha)$, and so can be extended to a basis $\alpha(v_1), \dots, \alpha(v_r), x_1, \dots, x_t$ for W . Now we will use the bases

$$\begin{aligned} v_1, \dots, v_r, v_{r+1} &= u_1, \dots, v_{r+s} = w_s & \text{for } V, \\ w_1 = \alpha(v_1), \dots, w_r = \alpha(v_r), w_{r+1} &= x_1, \dots, w_{r+s} = x_s & \text{for } W. \end{aligned}$$

We have

$$\alpha(v_i) = \begin{cases} w_i & \text{if } 1 \leq i \leq r, \\ 0 & \text{otherwise;} \end{cases}$$

so the matrix of α relative to these bases is

$$\begin{bmatrix} I_r & O \\ O & O \end{bmatrix}$$

as claimed.

Proposition 4.3 Let α be a linear map on V which is represented by the matrix A relative to a basis B , and by the matrix A' relative to a basis B' . Let $P = P_{B,B'}$ be the transition matrix between the two bases. Then

$$A' = P^{-1}AP.$$

Proof This is just Proposition 4.6, since P and Q are the same here.

Definition 4.2 Two $n \times n$ matrices A and B are said to be *similar* if $B = P^{-1}AP$ for some invertible matrix P .

Thus similarity is an equivalence relation, and

two matrices are similar if and only if they represent the same linear map with respect to different bases.

There is no simple canonical form for similarity like the one for equivalence that we met earlier. For the rest of this section we look at a special class of matrices or linear maps, the “diagonalisable” ones, where we do have a nice simple representative of the similarity class. In the final section we give without proof a general result for the complex numbers.

Preview from Notesale.co.uk
page 49 of 124

4.3 Eigenvalues and eigenvectors

Definition 4.3 Let α be a linear map on V . A vector $v \in V$ is said to be an *eigenvector* of α , with *eigenvalue* $\lambda \in \mathbb{K}$, if $v \neq 0$ and $\alpha(v) = \lambda v$. The set $\{v : \alpha(v) = \lambda v\}$ consisting of the zero vector and the eigenvectors with eigenvalue λ is called the λ -*eigenspace* of α .

Note that we require that $v \neq 0$; otherwise the zero vector would be an eigenvector for any value of λ . With this requirement, each eigenvector has a unique eigenvalue: for if $\alpha(v) = \lambda v = \mu v$, then $(\lambda - \mu)v = 0$, and so (since $v \neq 0$) we have $\lambda = \mu$.

The name *eigenvalue* is a mixture of German and English; it means “characteristic value” or “proper value” (here “proper” is used in the sense of “property”). Another term used in older books is “latent root”. Here “latent” means “hidden”: the idea is that the eigenvalue is somehow hidden in a matrix representing α , and we have to extract it by some procedure. We’ll see how to do this soon.

for any vector $v \in V$; and, if $v, w \neq 0$, then we define the angle between them to be θ , where

$$\cos \theta = \frac{v \cdot w}{|v| \cdot |w|}.$$

For this definition to make sense, we need to know that

$$-|v| \cdot |w| \leq v \cdot w \leq |v| \cdot |w|$$

for any vectors v, w (since $\cos \theta$ lies between -1 and 1). This is the content of the *Cauchy–Schwarz inequality*:

Theorem 6.1 *If v, w are vectors in an inner product space then*

$$(v \cdot w)^2 \leq (v \cdot v)(w \cdot w).$$

Proof By definition, we have $(v + xw) \cdot (v + xw) \geq 0$ for any real number x . Expanding, we obtain

$$x^2(w \cdot w) + 2x(v \cdot w) + (v \cdot v) \geq 0.$$

This is a quadratic function in x . Since it is non-negative for all real x , either it has no real roots, or it has two equal real roots, thus its discriminant is non-positive, that is,

$$(v \cdot w)^2 - (v \cdot v)(w \cdot w) \leq 0,$$

as required.

There is essentially only one kind of inner product on a real vector space.

Definition 6.2 A basis $(v_1, \dots, v_n)$ for an inner product space is called *orthonormal* if $v_i \cdot v_j = \delta_{ij}$ (the Kronecker delta) for $1 \leq i, j \leq n$.

Remark: If vectors $v_1, \dots, v_n$ satisfy $v_i \cdot v_j = \delta_{ij}$, then they are necessarily linearly independent. For suppose that $c_1 v_1 + \dots + c_n v_n = 0$. Taking the inner product of this equation with v_i , we find that $c_i = 0$, for all i .

Theorem 6.2 *Let $\cdot$ be an inner product on a real vector space V . Then there is an orthonormal basis $(v_1, \dots, v_n)$ for V . If we represent vectors in coordinates with respect to this basis, say $v = [x_1 \ x_2 \ \dots \ x_n]^T$ and $w = [y_1 \ y_2 \ \dots \ y_n]^T$, then*

$$v \cdot w = x_1 y_1 + x_2 y_2 + \dots + x_n y_n.$$

Proof This follows from our reduction of quadratic forms in the last chapter. Since the inner product is bilinear, the function $q(v) = v \cdot v = |v|^2$ is a quadratic form, and so it can be reduced to the form

$$q = x_1^2 + \cdots + x_s^2 - x_{s+1}^2 - \cdots - x_{s+t}^2.$$

Now we must have $s = n$ and $t = 0$. For, if $t > 0$, then the $s + 1$ st basis vector v_{s+1} satisfies $v_{s+1} \cdot v_{s+1} = -1$; while if $s + t < n$, then the n th basis vector v_n satisfies $v_n \cdot v_n = 0$. Either of these would contradict the positive definiteness of V . Now we have

$$q(x_1, \dots, x_n) = x_1^2 + \cdots + x_n^2,$$

and by polarisation we find that

$$b((x_1, \dots, x_n), (y_1, \dots, y_n)) = x_1 y_1 + \cdots + x_n y_n,$$

as required.

However, it is possible to give a more direct proof of the theorem; this is important because it involves a constructive method for finding an orthonormal basis, known as the *Gram–Schmidt process*.

Let $w_1, \dots, w_n$ be any basis for V . The Gram–Schmidt process works as follows.

- Since $w_1 \neq 0$, we have $w_1 \cdot w_1 > 0$, that is, $|w_1| > 0$. Put $v_1 = w_1/|w_1|$; then $|v_1| = 1$, that is, $v_1 \cdot v_1 = 1$.
- For $i = 2, \dots, n$, let $w'_i = w_i - (v_1 \cdot w_i)v_1$. Then

$$v_1 \cdot w'_i = v_1 \cdot w_i - (v_1 \cdot w_i)(v_1 \cdot v_1) = 0$$

for $i = 2, \dots, n$.

- Now apply the Gram–Schmidt process recursively to $(w'_2, \dots, w'_n)$.

Since we replace these vectors by linear combinations of themselves, their inner products with v_1 remain zero throughout the process. So if we end up with vectors $v_2, \dots, v_n$, then $v_1 \cdot v_i = 0$ for $i = 2, \dots, n$. By induction, we can assume that $v_i \cdot v_j = \delta_{ij}$ for $i, j = 2, \dots, n$; by what we have said, this holds if i or j is 1 as well.

Definition 6.3 The inner product on $\mathbb{R}^n$ for which the standard basis is orthonormal (that is, the one given in the theorem) is called the *standard inner product* on $\mathbb{R}^n$.

U , and so lie in $U^\perp$; and they are clearly linearly independent. Now suppose that $w \in U^\perp$ and $w = \sum c_i v_i$, where $(v_1, \dots, v_n)$ is the orthonormal basis we constructed. Then $c_i = w \cdot v_i = 0$ for $i = 1, \dots, r$; so w is a linear combination of the last $n - r$ basis vectors, which thus form a basis of $U^\perp$. Hence $\dim(U^\perp) = n - r$, as required.

Now the last statement of the proposition follows from the proof, since we have a basis for V which is a disjoint union of bases for U and $U^\perp$.

Recall the connection between direct sum decompositions and projections. If we have projections $P_1, \dots, P_r$ whose sum is the identity and which satisfy $P_i P_j = 0$ for $i \neq j$, then the space V is the direct sum of their images. This can be refined in an inner product space as follows.

Definition 7.2 Let V be an inner product space. A linear map $\pi : V \rightarrow V$ is an *orthogonal projection* if

- (a) π is a projection, that is, $\pi^2 = \pi$;
- (b) π is self-adjoint, that is, $\pi^* = \pi$ (where $\pi^*(v) \cdot w = v \cdot \pi(w)$ for all $v, w \in V$).

Proposition 7.2 If π is an orthogonal projection, then $\text{Ker}(\pi) = \text{Im}(\pi)^\perp$.

Proof We know that $V = \text{Ker}(\pi) \oplus \text{Im}(\pi)$; we only have to show that these two subspaces are orthogonal. So take $v \in \text{Ker}(\pi)$, so that $\pi(v) = 0$, and $w \in \text{Im}(\pi)$, so that $w = \pi(u)$ for some $u \in V$. Then

$$v \cdot w = v \cdot \pi(u) = \pi^*(v) \cdot u = \pi(v) \cdot u = 0,$$

as required.

Proposition 7.3 Let $\pi_1, \dots, \pi_r$ be orthogonal projections on an inner product space V satisfying $\pi_1 + \dots + \pi_r = I$ and $\pi_i \pi_j = 0$ for $i \neq j$. Let $U_i = \text{Im}(\pi_i)$ for $i = 1, \dots, r$. Then

$$V = U_1 \oplus U_2 \oplus \dots \oplus U_r,$$

and if $u_i \in U_i$ and $u_j \in U_j$, then u_i and u_j are orthogonal.

Proof The fact that V is the direct sum of the images of the π_i follows from Proposition 5.2. We only have to prove the last part. So take u_i and u_j as in the Proposition, say $u_i = \pi_i(v)$ and $u_j = \pi_j(w)$. Then

$$u_i \cdot u_j = \pi_i(v) \cdot \pi_j(w) = \pi_i^*(v) \cdot \pi_j(w) = v \cdot \pi_i(\pi_j(w)) = 0,$$

where the second equality holds since π_i is self-adjoint and the third is the definition of the adjoint.

A direct sum decomposition satisfying the conditions of the theorem is called an *orthogonal decomposition* of V .

Conversely, if we are given an orthogonal decomposition of V , then we can find orthogonal projections satisfying the hypotheses of the theorem.

Proof This is obvious since if b is alternating then $a_{ji} = b(v_j, v_i) = -b(v_i, v_j) = -a_{ij}$ and $a_{ii} = b(v_i, v_i) = 0$.

So we can write our theorem in matrix form as follows:

Theorem 9.4 Let A be an alternating matrix (or a skew-symmetric matrix over a field whose characteristic is not equal to 2). Then there is an invertible matrix P such that $P^\top AP$ is the matrix with s blocks $\begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix}$ on the diagonal and all other entries zero. Moreover the number s is half the rank of A , and so is independent of the choice of P .

Proof We know that the effect of a change of basis with transition matrix P is to replace the matrix A representing a bilinear form by $P^\top AP$. Also, the matrix in the statement of the theorem is just the matrix representing b relative to the special basis that we found in the preceding theorem.

This has a corollary which is a bit surprising at first sight:

Corollary 9.5 (a) The rank of a skew-symmetric matrix (over a field of characteristic not equal to 2) is even.

(b) The determinant of a skew-symmetric matrix (over a field of characteristic not equal to 2) is a square, and is zero if the size of the matrix is odd.

Proof (a) The canonical form in the theorem clearly has rank $2s$.

(b) If the skew-symmetric matrix A is singular then its determinant is zero, which is a square. So suppose that it is invertible. Then its canonical form has $s = n/2$ blocks $\begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix}$ on the diagonal. Each of these blocks has determinant 1, and hence so does the whole matrix. So $\det(P^\top AP) = \det(P)^2 \det(A) = 1$, whence $\det(A) = 1/(\det(P)^2)$, which is a square.

If the size n of A is odd, then the rank cannot be n (by (a)), and so $\det(A) = 0$.

Remark There is a function defined on skew-symmetric matrices called the *Pfaffian*, which like the determinant is a polynomial in the matrix entries, and has the property that $\det(A)$ is the square of the Pfaffian of A : that is, $\det(A) = (\text{Pf}(A))^2$.

For example,

$$\text{Pf} \begin{bmatrix} 0 & a \\ -a & 0 \end{bmatrix} = a, \quad \text{Pf} \begin{bmatrix} 0 & a & b & c \\ -a & 0 & d & e \\ -b & -d & 0 & f \\ -c & -e & -f & 0 \end{bmatrix} = af - be + cd.$$

(Check that the determinant of the second matrix is $(af - be + cd)^2$.)

Preview from Notesale.co.uk
Page 109 of 124

Appendix F

Worked examples

1. Let

$$A = \begin{bmatrix} 1 & 2 & 4 & -1 & 5 \\ 1 & 2 & 3 & -1 & 3 \\ -1 & -2 & 0 & 1 & 3 \end{bmatrix}.$$

- (a) Find a basis for the row space of A .
- (b) What is the rank of A ?
- (c) Find a basis for the column space of A .
- (d) Find invertible matrices P and Q such that $P^{-1}AQ$ is in the canonical form for equivalence.

(a) Subtract the first row from the second, add the first row to the third, then multiply the new second row by -1 and subtract four times this row from the third, to get the matrix

$$B = \begin{bmatrix} 1 & 2 & 4 & -1 & 5 \\ 0 & 0 & 1 & 0 & 2 \\ 0 & 0 & 0 & 0 & 0 \end{bmatrix}.$$

The first two rows clearly form a basis for the row space.

(b) The rank is 2, since there is a basis with two elements.

(c) The column rank is equal to the row rank and so is also equal to 2. By inspection, the first and third columns of A are linearly independent, so they form a basis. The first and second columns are not linearly independent, so we cannot use these! (Note that we have to go back to the original A here; row operations change the column space, so selecting two independent columns of B would not be correct.)

Remark A more elegant solution is the matrix

$$\frac{1}{2} \begin{bmatrix} 1 & 1 & 1 & 1 \\ 1 & -1 & 1 & -1 \\ 1 & 1 & -1 & -1 \\ 1 & -1 & -1 & 1 \end{bmatrix}.$$

This matrix (without the factor $\frac{1}{2}$) is known as a *Hadamard matrix*. It is an $n \times n$ matrix H with all entries ± 1 satisfying $H^T H = nI$. It is known that an $n \times n$ Hadamard matrix cannot exist unless n is 1, 2, or a multiple of 4; however, nobody has succeeded in proving that a Hadamard matrix of any size n divisible by 4 exists.

The smallest order for which the existence of a Hadamard matrix is still in doubt is (at the time of writing) $n = 668$. The previous smallest, $n = 428$, was resolved only in 2004 by Hadi Kharaghani and Behruz Tayfeh-Rezaie in Tehran, by constructing an example.

As a further exercise, show that, if H is a Hadamard matrix of size n , then $\begin{bmatrix} H & H \\ H & -H \end{bmatrix}$ is a Hadamard matrix of size $2n$. (The Hadamard matrix of size 4 constructed above is of this form.)

8. Let $A = \begin{bmatrix} 1 & 1 \\ 1 & 2 \end{bmatrix}$ and $B = \begin{bmatrix} 1 & 1 \\ 1 & 0 \end{bmatrix}$.

Find an invertible matrix P and a diagonal matrix D such that $P^T A P = I$ and $P^T B P = D$, where I is the identity matrix.

First we take the quadratic form corresponding to A , and reduce it to a sum of squares. The form is $x^2 + 2xy + 2y^2$, which is $(x+y)^2 + y^2$. (Note: This is the sum of two squares, in agreement with the fact that A is positive definite.)

Now the matrix that transforms (x, y) to $(x+y, y)$ is $Q = \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix}$, since

$$\begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} = \begin{bmatrix} x+y \\ y \end{bmatrix}.$$

Hence

$$\begin{bmatrix} x & y \end{bmatrix} Q^T Q \begin{bmatrix} x \\ y \end{bmatrix} = x^2 + 2xy + 2y^2 = \begin{bmatrix} x & y \end{bmatrix} A \begin{bmatrix} x \\ y \end{bmatrix},$$

so that $Q^T Q = A$.

Now, if we put $P = Q^{-1} = \begin{bmatrix} 1 & -1 \\ 0 & 1 \end{bmatrix}$, we see that $P^T A P = P^T (Q^T Q) P = I$.