

Appendix A. Linear Algebra

Sciavicco L., Siciliano B.: Modelling and Control of Robot Manipulators, Edition 2, Springer Verlag, London, 2005, 6th printing, pp. 335-349

Since modelling and control of robot manipulators requires an extensive use of *matrices* and *vectors* as well as of matrix and vector *operations*, the goal of this appendix is to provide a brush-up of *linear algebra*.

A.1 Definitions

A *matrix* of dimensions $(m \times n)$, with m and n positive integers, is an array of elements a_{ij} arranged into m rows and n columns:

$$A = [a_{ij}]_{\substack{i=1,\dots,m \\ j=1,\dots,n}} = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{bmatrix}. \quad (\text{A.1})$$

If $m = n$, the matrix is said to be *square*; if $m < n$, the matrix has more columns than rows; if $m > n$ the matrix has more rows than columns. Further, if $n = 1$, the notation (A.1) is used to represent a (column) vector $\mathbf{a}$ of dimensions $(m \times 1)^1$; the elements a_i are said to be vector components. A square matrix $\mathbf{A}$ of dimensions $(n \times n)$ is said to be *upper triangular* if $a_{ij} = 0$ for $i > j$:

$$\mathbf{A} = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ 0 & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & a_{nn} \end{bmatrix};$$

the matrix is said to be *lower triangular* if $a_{ij} = 0$ for $i < j$.

An $(n \times n)$ square matrix $\mathbf{A}$ is said to be *diagonal* if $a_{ij} = 0$ for $i \neq j$, i.e.,

$$\mathbf{A} = \begin{bmatrix} a_{11} & 0 & \dots & 0 \\ 0 & a_{22} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & a_{nn} \end{bmatrix} = \text{diag}\{a_{11}, a_{22}, \dots, a_{nn}\}.$$

¹ According to standard mathematical notation, small boldface is used to denote vectors while capital boldface is used to denote matrices. Scalars are denoted by roman characters.

If an $(n \times n)$ diagonal matrix has all unit elements on the diagonal ($a_{ii} = 1$), the matrix is said to be *identity* and is denoted by I_n^2 . A matrix is said to be *null* if all its elements are null and is denoted by O . The null column vector is denoted by 0 .

The *transpose* A^T of a matrix A of dimensions $(m \times n)$ is the matrix of dimensions $(n \times m)$ which is obtained from the original matrix by interchanging its rows and columns:

$$A^T = \begin{bmatrix} a_{11} & a_{21} & \dots & a_{m1} \\ a_{12} & a_{22} & \dots & a_{m2} \\ \vdots & \vdots & \ddots & \vdots \\ a_{1n} & a_{2n} & \dots & a_{mn} \end{bmatrix}. \quad (\text{A.2})$$

The transpose of a column vector a is the row vector a^T .

An $(n \times n)$ square matrix A is said to be *symmetric* if $A^T = A$, and thus $a_{ij} = a_{ji}$:

$$A = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{12} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{1n} & a_{2n} & \dots & a_{nn} \end{bmatrix}.$$

An $(n \times n)$ square matrix A is said to be *skew-symmetric* if $A^T = -A$, and thus $a_{ij} = -a_{ji}$ for $i \neq j$ and $a_{ii} = 0$, leading to

$$A = \begin{bmatrix} 0 & a_{12} & \dots & a_{1n} \\ -a_{12} & 0 & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ -a_{1n} & -a_{2n} & \dots & 0 \end{bmatrix}.$$

A *partitioned* matrix is a matrix whose elements are matrices (*blocks*) of proper dimensions:

$$A = \begin{bmatrix} A_{11} & A_{12} & \dots & A_{1n} \\ A_{21} & A_{22} & \dots & A_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ A_{m1} & A_{m2} & \dots & A_{mn} \end{bmatrix}.$$

A partitioned matrix may be block-triangular or block-diagonal. Special partitions of a matrix are that by columns

$$A = [a_1 \quad a_2 \quad \dots \quad a_n]$$

and that by rows

$$A = \begin{bmatrix} a_1^T \\ a_2^T \\ \vdots \\ a_m^T \end{bmatrix}.$$

² Subscript n is usually omitted if the dimensions are clear from the context.

Given a square matrix A of dimensions $(n \times n)$, the *algebraic complement* $A_{(ij)}$ of element a_{ij} is the matrix of dimensions $((n-1) \times (n-1))$ which is obtained by eliminating row i and column j of matrix A .

A.2 Matrix Operations

The *trace* of an $(n \times n)$ square matrix A is the sum of the elements on the diagonal:

$$\text{Tr}(A) = \sum_{i=1}^n a_{ii}. \quad (\text{A.3})$$

Two matrices A and B of the same dimensions $(m \times n)$ are equal if $a_{ij} = b_{ij}$. If A and B are two matrices of the same dimensions, their *sum* is the matrix

$$C = A + B \quad (\text{A.4})$$

whose elements are given by $c_{ij} = a_{ij} + b_{ij}$. The following properties hold:

$$\begin{aligned} A + O &= A \\ A + B &= B + A \\ (A + B) + C &= A + (B + C). \end{aligned}$$

Notice that two matrices of the same dimensions and partitioned in the same way can be summed formally by operating on the blocks in the same position and treating them like elements.

The product of a scalar α by an $(m \times n)$ matrix A is the matrix αA whose elements are given by αa_{ij} . If A is an $(n \times n)$ diagonal matrix with all equal elements on the diagonal ($a_{ii} = a$), it follows that $A = aI_n$.

If A is a square matrix, one may write

$$A = A_s + A_a \quad (\text{A.5})$$

where

$$A_s = \frac{1}{2}(A + A^T) \quad (\text{A.6})$$

is a symmetric matrix representing the *symmetric* part of A , and

$$A_a = \frac{1}{2}(A - A^T) \quad (\text{A.7})$$

is a skew-symmetric matrix representing the *skew-symmetric* part of A .

The row-by-column *product* of a matrix A of dimensions $(m \times p)$ by a matrix B of dimensions $(p \times n)$ is the matrix of dimensions $(m \times n)$

$$C = AB \quad (\text{A.8})$$

whose elements are given by $c_{ij} = \sum_{k=1}^p a_{ik}b_{kj}$. The following properties hold:

$$\begin{aligned} A &= AI_p = I_m A \\ A(BC) &= (AB)C \\ A(B+C) &= AB + AC \\ (A+B)C &= AC + BC \\ (AB)^T &= B^T A^T. \end{aligned}$$

Notice that, in general, $AB \neq BA$, and $AB = O$ does not imply that $A = O$ or $B = O$; further, notice that $AC = BC$ does not imply that $A = B$.

If an $(m \times p)$ matrix A and a $(p \times n)$ matrix B are partitioned in such a way that the number of blocks for each row of A is equal to the number of blocks for each column of B , and the blocks A_{ik} and B_{kj} have dimensions compatible with product, the matrix product AB can be formally obtained by operating by rows and columns on the blocks of proper position and treating them like elements.

For an $(n \times n)$ square matrix A , the *determinant* of A is the scalar given by the following expression, which holds $\forall i = 1, \dots, n$:

$$\det(A) = \sum_{j=1}^n a_{ij}(-1)^{i+j} \det(A_{(ij)}). \quad (\text{A.9})$$

The determinant can be computed according to any row i as in (A.9); the same result is obtained by computing it according to any column j . If $n = 1$, then $\det(a_{11}) = a_{11}$. The following property holds:

$$\det(A) = \det(A^T).$$

Moreover, interchanging two generic columns p and q of a matrix A yields

$$\det([a_1 \dots a_p \dots a_q \dots a_n]) = -\det([a_1 \dots a_q \dots a_p \dots a_n]).$$

As a consequence, if a matrix has two equal columns (rows), then its determinant is null. Also, it is $\det(\alpha A) = \alpha^n \det(A)$.

Given an $(m \times n)$ matrix A , the determinant of the square block obtained by selecting an equal number k of rows and columns is said to be k -order *minor* of matrix A . The minors obtained by taking the *first* k rows and columns of A are said to be *principal* minors.

If A and B are square matrices, then

$$\det(AB) = \det(A)\det(B). \quad (\text{A.10})$$

If A is an $(n \times n)$ triangular matrix (in particular diagonal), then

$$\det(A) = \prod_{i=1}^n a_{ii}.$$

More generally, if A is block-triangular with m blocks A_{ii} on the diagonal, then

$$\det(A) = \prod_{i=1}^m \det(A_{ii}).$$

A square matrix A is said to be *singular* when $\det(A) = 0$.

The *rank* $\rho(A)$ of a matrix A of dimensions $(m \times n)$ is the maximum integer r so that at least a nonnull minor of order r exists. The following properties hold:

$$\begin{aligned}\rho(A) &\leq \min\{m, n\} \\ \rho(A) &= \rho(A^T) \\ \rho(A^T A) &= \rho(A) \\ \rho(AB) &\leq \min\{\rho(A), \rho(B)\}.\end{aligned}$$

A matrix so that $\rho(A) = \min\{m, n\}$ is said to be *full-rank*.

The *adjoint* of a square matrix A is the matrix

$$\text{Adj } A = [(-1)^{i+j} \det(A_{(ij)})]_{i=1, \dots, n}^T \quad (A.11)$$

An $(n \times n)$ square matrix A is said to be *invertible* if a matrix A^{-1} exists, termed *inverse* of A , so that

$$A^{-1}A = AA^{-1} = I_n.$$

Since $\rho(I_n) = n$, an $(n \times n)$ square matrix A is invertible if and only if $\rho(A) = n$, i.e., $\det(A) \neq 0$ (nonsingular matrix). The inverse of A can be computed as

$$A^{-1} = \frac{1}{\det(A)} \text{Adj } A. \quad (A.12)$$

The following properties hold:

$$\begin{aligned}(A^{-1})^{-1} &= A \\ (A^T)^{-1} &= (A^{-1})^T.\end{aligned}$$

If the inverse of a square matrix is equal to its transpose

$$A^T = A^{-1}, \quad (A.13)$$

then the matrix is said to be *orthogonal*; in this case it is

$$AA^T = A^T A = I. \quad (A.14)$$

If A and B are invertible square matrices of the same dimensions, then

$$(AB)^{-1} = B^{-1}A^{-1}. \quad (A.15)$$

Given n square matrices A_{ii} all invertible, the following expression holds:

$$(\text{diag}\{A_{11}, \dots, A_{nn}\})^{-1} = \text{diag}\{A_{11}^{-1}, \dots, A_{nn}^{-1}\}.$$

where $\text{diag}\{A_{11}, \dots, A_{nn}\}$ denotes the block-diagonal matrix.

If A and C are invertible square matrices of proper dimensions, the following expression holds:

$$(A + BCD)^{-1} = A^{-1} - A^{-1}B(DA^{-1}B + C^{-1})^{-1}DA^{-1},$$

where the matrix $DA^{-1}B + C^{-1}$ must be invertible.

If a block-partitioned matrix is invertible, then its inverse is given by the general expression

$$\begin{bmatrix} A & D \\ C & B \end{bmatrix}^{-1} = \begin{bmatrix} A^{-1} + E\Delta^{-1}F & -E\Delta^{-1} \\ -\Delta^{-1}F & \Delta^{-1} \end{bmatrix} \quad (\text{A.16})$$

where $\Delta = B - CA^{-1}D$, $E = A^{-1}D$ and $F = CA^{-1}$, on the assumption that the inverses of matrices A and Δ exist. In the case of a block-triangular matrix, invertibility of the matrix requires invertibility of the blocks on the diagonal. The following expressions hold:

$$\begin{aligned} \begin{bmatrix} A & O \\ C & B \end{bmatrix}^{-1} &= \begin{bmatrix} A^{-1} & O \\ -B^{-1}CA^{-1} & B^{-1} \end{bmatrix} \\ \begin{bmatrix} A & D \\ O & B \end{bmatrix}^{-1} &= \begin{bmatrix} A^{-1} & -A^{-1}DB^{-1} \\ O & B^{-1} \end{bmatrix}. \end{aligned}$$

The derivative of an $(m \times n)$ matrix $A(t)$, whose elements $a_{ij}(t)$ are differentiable functions, is the matrix

$$\dot{A}(t) = \frac{d}{dt}A(t) = \left[\frac{d}{dt}a_{ij}(t) \right]_{\substack{i=1, \dots, m \\ j=1, \dots, n}}. \quad (\text{A.17})$$

If an $(n \times n)$ square matrix $A(t)$ is so that $\rho(A(t)) = n \forall t$ and its elements $a_{ij}(t)$ are differentiable functions, then the derivative of the inverse of $A(t)$ is given by

$$\frac{d}{dt}A^{-1}(t) = -A^{-1}(t)\dot{A}(t)A^{-1}(t). \quad (\text{A.18})$$

Given a scalar function $f(x)$, endowed with partial derivatives with respect to the elements x_i of the $(n \times 1)$ vector x , the gradient of function f with respect to vector x is the $(n \times 1)$ column vector

$$\text{grad}_x f(x) = \left(\frac{\partial f(x)}{\partial x} \right)^T = \left[\frac{\partial f(x)}{\partial x_1} \quad \frac{\partial f(x)}{\partial x_2} \quad \dots \quad \frac{\partial f(x)}{\partial x_n} \right]^T. \quad (\text{A.19})$$

Further, if $x(t)$ is a differentiable function with respect to t , then

$$\dot{f}(x) = \frac{d}{dt}f(x(t)) = \frac{\partial f}{\partial x} \dot{x} = \text{grad}_x^T f(x) \dot{x}. \quad (\text{A.20})$$

Given a vector function $g(x)$ of dimensions $(m \times 1)$, whose elements g_i are differentiable with respect to the vector x of dimensions $(n \times 1)$, the Jacobian matrix (or simply *Jacobian*) of the function is defined as the $(m \times n)$ matrix

$$J_g(x) = \frac{\partial g(x)}{\partial x} = \begin{bmatrix} \frac{\partial g_1(x)}{\partial x} \\ \frac{\partial g_2(x)}{\partial x} \\ \vdots \\ \frac{\partial g_m(x)}{\partial x} \end{bmatrix}. \quad (\text{A.21})$$

If $x(t)$ is a differentiable function with respect to t , then

$$\dot{g}(x) = \frac{d}{dt}g(x(t)) = \frac{\partial g}{\partial x} \dot{x} = J_g(x) \dot{x}. \quad (\text{A.22})$$

A.3 Vector Operations

Given n vectors x_i of dimensions $(m \times 1)$, they are said to be *linearly independent* if the expression

$$k_1 x_1 + k_2 x_2 + \dots + k_n x_n = 0$$

holds only when all the constants k_i vanish. A necessary and sufficient condition for the vectors $x_1, x_2, \dots, x_n$ to be linearly independent is that the matrix

$$A = [x_1 \quad x_2 \quad \dots \quad x_n]$$

has rank n ; this implies that a necessary condition for linear independence is that $n \leq m$. If instead $\rho(A) = r < n$, then only r vectors are linearly independent and the remaining $n - r$ vectors can be expressed as a linear combination of the previous ones.

A system of vectors $\mathcal{X}$ is a *vector space* on the field of real numbers $\mathbb{R}$ if the operations of *sum of two vectors* of $\mathcal{X}$ and *product of a scalar by a vector* of $\mathcal{X}$ have values in $\mathcal{X}$ and the following properties hold:

$$\begin{aligned}
x + y &= y + x & \forall x, y \in \mathcal{X} \\
(x + y) + z &= x + (y + z) & \forall x, y, z \in \mathcal{X} \\
\exists 0 \in \mathcal{X} : x + 0 &= x & \forall x \in \mathcal{X} \\
\forall x \in \mathcal{X}, \exists (-x) \in \mathcal{X} : x + (-x) &= 0 \\
1x &= x & \forall x \in \mathcal{X} \\
\alpha(\beta x) &= (\alpha\beta)x & \forall \alpha, \beta \in \mathbb{R} \quad \forall x \in \mathcal{X} \\
(\alpha + \beta)x &= \alpha x + \beta x & \forall \alpha, \beta \in \mathbb{R} \quad \forall x \in \mathcal{X} \\
\alpha(x + y) &= \alpha x + \alpha y & \forall \alpha \in \mathbb{R} \quad \forall x, y \in \mathcal{X}.
\end{aligned}$$

The *dimension* of the space $\dim(\mathcal{X})$ is the maximum number of linearly independent vectors x in the space. A set $\{x_1, x_2, \dots, x_n\}$ of linearly independent vectors is a *basis* of vector space $\mathcal{X}$, and each vector y in the space can be uniquely expressed as a linear combination of vectors from the basis:

$$y = c_1 x_1 + c_2 x_2 + \dots + c_n x_n \quad (\text{A.23})$$

where the constants $c_1, c_2, \dots, c_n$ are said to be the *components* of the vector y in the basis $\{x_1, x_2, \dots, x_n\}$.

A subset $\mathcal{Y}$ of a vector space $\mathcal{X}$ is a *subspace* $\mathcal{Y} \subseteq \mathcal{X}$ if it is a vector space with the operations of vector sum and product of a scalar by a vector, *i.e.*,

$$\alpha x + \beta y \in \mathcal{Y} \quad \forall \alpha, \beta \in \mathbb{R} \quad \forall x, y \in \mathcal{Y}.$$

According to a geometric interpretation, a subspace is a *hyperplane* passing by the origin (null element) of $\mathcal{X}$.

The *scalar product* $\langle x, y \rangle$ of two vectors x and y of dimensions $(m \times 1)$ is the scalar that is obtained by summing the products of the respective components in a given basis:

$$\langle x, y \rangle = x_1 y_1 + x_2 y_2 + \dots + x_m y_m = x^T y = y^T x. \quad (\text{A.24})$$

Two vectors are said to be *orthogonal* when their scalar product is null:

$$x^T y = 0. \quad (\text{A.25})$$

The *norm* of a vector can be defined as

$$\|x\| = \sqrt{x^T x}. \quad (\text{A.26})$$

It is possible to show that both the *triangle inequality*

$$\|x + y\| \leq \|x\| + \|y\| \quad (\text{A.27})$$

and the *Schwarz' inequality*

$$|x^T y| \leq \|x\| \|y\| \quad (\text{A.28})$$

hold. A *unit vector* $\hat{x}$ is a vector whose *norm* is unity, i.e., $\hat{x}^T \hat{x} = 1$. Given a vector x , its unit vector is obtained by dividing each component by its norm:

$$\hat{x} = \frac{1}{\|x\|} x. \quad (\text{A.29})$$

A typical example of vector space is the *Euclidean space* whose dimension is 3; in this case a basis is constituted by the unit vectors of a coordinate frame.

The *vector product* of two vectors x and y in the Euclidean space is the vector

$$x \times y = \begin{bmatrix} x_2 y_3 - x_3 y_2 \\ x_3 y_1 - x_1 y_3 \\ x_1 y_2 - x_2 y_1 \end{bmatrix}. \quad (\text{A.30})$$

The following properties hold:

$$\begin{aligned} x \times x &= 0 \\ x \times y &= -y \times x \\ x \times (y + z) &= x \times y + x \times z. \end{aligned}$$

The vector product of two vectors x and y can be expressed also as the product of a matrix operator $S(x)$ by the vector y . In fact, by introducing the *skew-symmetric* matrix

$$S(x) = \begin{bmatrix} 0 & -x_3 & x_2 \\ x_3 & 0 & -x_1 \\ -x_2 & x_1 & 0 \end{bmatrix} \quad (\text{A.31})$$

obtained with the components of vector x , the vector product $x \times y$ is given by

$$x \times y = S(x)y = -S(y)x \quad (\text{A.32})$$

as can be easily verified. Moreover, the following properties hold:

$$\begin{aligned} S(x)x &= S^T(x)x = 0 \\ S(\alpha x + \beta y) &= \alpha S(x) + \beta S(y) \\ S(x)S(y) &= yx^T - x^T y I \\ S(S(x)y) &= yx^T - xy^T. \end{aligned}$$

Given three vector x, y, z in the Euclidean space, the following expressions hold for the *scalar triple products*:

$$x^T(y \times z) = y^T(z \times x) = z^T(x \times y). \quad (\text{A.33})$$

If any two vectors of three are equal, then the scalar triple product is null; e.g.,

$$x^T(x \times y) = 0.$$

A.4 Linear Transformations

Consider a vector space $\mathcal{X}$ of dimension n and a vector space $\mathcal{Y}$ of dimension m with $m \leq n$. The *linear transformation* between the vectors $x \in \mathcal{X}$ and $y \in \mathcal{Y}$ can be defined as

$$y = Ax \quad (\text{A.34})$$

in terms of the matrix A of dimensions $(m \times n)$. The *range space* (or simply *range*) of the transformation is the subspace

$$\mathcal{R}(A) = \{y : y = Ax, x \in \mathcal{X}\} \subseteq \mathcal{Y}, \quad (\text{A.35})$$

which is the subspace generated by the linearly independent columns of matrix A taken as a basis of $\mathcal{Y}$. It is easy to recognize that

$$\varrho(A) = \dim(\mathcal{R}(A)). \quad (\text{A.36})$$

On the other hand, the *null space* (or simply *null*) of the transformation is the subspace

$$\mathcal{N}(A) = \{x : Ax = 0, x \in \mathcal{X}\} \subseteq \mathcal{X}. \quad (\text{A.37})$$

Given a matrix A of dimensions $(m \times n)$, the notable result holds:

$$\varrho(A) + \dim(\mathcal{N}(A)) = n. \quad (\text{A.38})$$

Therefore, if $\varrho(A) = r \leq \min\{m, n\}$, then $\dim(\mathcal{R}(A)) = r$ and $\dim(\mathcal{N}(A)) = n - r$. It follows that if $m < n$, then $\mathcal{N}(A) \neq \emptyset$ independently of the rank of A ; if $m = n$, then $\mathcal{N}(A) \neq \emptyset$ only in the case of $\varrho(A) = r < m$.

If $x \in \mathcal{N}(A)$ and $y \in \mathcal{R}(A^T)$, then $y^T x = 0$, i.e., the vectors in the null space of A are orthogonal to each vector in the range space of the transpose of A . It can be shown that the set of vectors orthogonal to each vector of the range space of A^T coincides with the null space of A , whereas the set of vectors orthogonal to each vector in the null space of A^T coincides with the range space of A . In symbols:

$$\mathcal{N}(A) \equiv \mathcal{R}^\perp(A^T) \quad \mathcal{R}(A) \equiv \mathcal{N}^\perp(A^T) \quad (\text{A.39})$$

where $\perp$ denotes the *orthogonal complement* of a subspace.

A linear transformation allows defining the *norm* of a matrix A induced by the norm defined for a vector x as follows. In view of the property

$$\|Ax\| \leq \|A\| \|x\|, \quad (\text{A.40})$$

the norm of A can be defined as

$$\|A\| = \sup_{x \neq 0} \frac{\|Ax\|}{\|x\|} \quad (\text{A.41})$$

which can be computed also as

$$\max_{\|x\|=1} \|Ax\|.$$

A direct consequence of (A.40) is the property

$$\|AB\| \leq \|A\| \|B\|. \quad (\text{A.42})$$

A.5 Eigenvalues and Eigenvectors

Consider the linear transformation on a vector u established by an $(n \times n)$ square matrix A . If the vector resulting from the transformation has the same direction of u (with $u \neq 0$), then

$$Au = \lambda u. \quad (\text{A.43})$$

The equation in (A.43) can be rewritten in matrix form as

$$(\lambda I - A)u = 0. \quad (\text{A.44})$$

For the homogeneous system of equations in (A.44) to have a solution different from the trivial one $u = 0$, it must be

$$\det(\lambda I - A) = 0 \quad (\text{A.45})$$

which is termed *characteristic equation*. Its solutions $\lambda_1, \dots, \lambda_n$ are the *eigenvalues* of matrix A ; they coincide with the eigenvalues of matrix A^T . On the assumption of distinct eigenvalues, the n vectors u_i satisfying the equation

$$(\lambda_i I - A)u_i = 0 \quad i = 1, \dots, n \quad (\text{A.46})$$

are said to be the *eigenvectors* associated with the eigenvalues λ_i .

The matrix U formed by the column vectors u_i is invertible and constitutes a basis in the space of dimension n . Further, the *similarity transformation* established by U :

$$A = U^{-1}AU \quad (\text{A.47})$$

is so that $A = \text{diag}\{\lambda_1, \dots, \lambda_n\}$. It follows that $\det(A) = \prod_{i=1}^n \lambda_i$.

If the matrix A is *symmetric*, its eigenvalues are real and A can be written as

$$A = U^T AU; \quad (\text{A.48})$$

hence, the eigenvector matrix U is orthogonal.

A.6 Bilinear Forms and Quadratic Forms

A *bilinear form* in the variables x_i and y_j is the scalar

$$B = \sum_{i=1}^m \sum_{j=1}^n a_{ij} x_i y_j$$

which can be written in matrix form

$$B(x, y) = x^T A y = y^T A^T x \quad (\text{A.49})$$

where $x = [x_1 \ x_2 \ \dots \ x_m]^T$, $y = [y_1 \ y_2 \ \dots \ y_n]^T$, and A is the $(m \times n)$ matrix of the coefficients a_{ij} representing the core of the form.

A special case of bilinear form is the *quadratic form*

$$Q(x) = x^T A x \quad (\text{A.50})$$

where A is an $(n \times n)$ square matrix. Hence, for computation of (A.50), the matrix A can be replaced with its symmetric part A_s given by (A.6). It follows that if A is a *skew-symmetric* matrix, then

$$x^T A x = 0 \quad \forall x.$$

The quadratic form (A.50) is said to be *positive definite* if

$$\begin{aligned} x^T A x &> 0 & \forall x \neq 0 \\ x^T A x &= 0 & x = 0. \end{aligned} \quad (\text{A.51})$$

The matrix A core of the form is also said to be *positive definite*. Analogously, a quadratic form is said to be *negative definite* if it can be written as $-Q(x) = -x^T A x$ where $Q(x)$ is positive definite.

A necessary condition for a square matrix to be positive definite is that its elements on the diagonal are strictly positive. Further, in view of (A.48), the eigenvalues of a positive definite matrix are all positive. If the eigenvalues are not known, a necessary and sufficient condition for a symmetric matrix to be positive definite is that its principal minors are strictly positive (*Sylvester's criterion*). It follows that a positive definite matrix is full-rank and thus it is always invertible.

A symmetric positive definite matrix A can always be decomposed as

$$A = U^T \Lambda U, \quad (\text{A.52})$$

where U is an orthogonal matrix of eigenvectors ($U^T U = I$) and Λ is the diagonal matrix of the eigenvalues of A .

Let $\lambda_{\min}(A)$ and $\lambda_{\max}(A)$ respectively denote the smallest and largest eigenvalues of a positive definite matrix A ($\lambda_{\min}, \lambda_{\max} > 0$). Then, the quadratic form in (A.50) satisfies the following inequality:

$$\lambda_{\min}(A) \|x\|^2 \leq x^T A x \leq \lambda_{\max}(A) \|x\|^2. \quad (\text{A.53})$$

An $(n \times n)$ square matrix A is said to be *positive semi-definite* if

$$x^T A x \geq 0 \quad \forall x. \quad (\text{A.54})$$

This definition implies that $\rho(A) = r < n$, and thus r eigenvalues of A are positive and $n - r$ are null. Therefore, a positive semi-definite matrix A has a null space of finite

dimension, and specifically the form vanishes when $x \in \mathcal{N}(A)$. A typical example of a positive semi-definite matrix is the matrix $A = H^T H$ where H is an $(m \times n)$ matrix with $m < n$. In an analogous way, a *negative semi-definite* matrix can be defined.

Given the *bilinear form* in (A.49), the *gradient* of the form with respect to x is given by

$$\text{grad}_x B(x, y) = \left(\frac{\partial B(x, y)}{\partial x} \right)^T = Ay, \quad (\text{A.55})$$

whereas the gradient of B with respect to y is given by

$$\text{grad}_y B(x, y) = \left(\frac{\partial B(x, y)}{\partial y} \right)^T = A^T x. \quad (\text{A.56})$$

Given the *quadratic form* in (A.50) with A *symmetric*, the *gradient* of the form with respect to x is given by

$$\text{grad}_x Q(x) = \left(\frac{\partial Q(x)}{\partial x} \right)^T = 2Ax. \quad (\text{A.57})$$

Further, if x and A are differentiable functions of t , then

$$\dot{Q}(x) = \frac{d}{dt} Q(x(t)) = 2x^T A \dot{x} + x^T \dot{A} x; \quad (\text{A.58})$$

if A is constant, then the second term obviously vanishes.

A.7 Pseudo-inverse

The inverse of a matrix can be defined only when the matrix is square and nonsingular. The inverse operation can be extended to the case of non-square matrices. Given a matrix A of dimensions $(m \times n)$ with $\rho(A) = \min\{m, n\}$, if $n < m$, a *left inverse* of A can be defined as the matrix A_l of dimensions $(n \times m)$ so that

$$A_l A = I_n.$$

If instead $n > m$, a *right inverse* of A can be defined as the matrix A_r of dimensions $(n \times m)$ so that

$$A A_r = I_m.$$

If A has more rows than columns ($m > n$) and has rank n , a special left inverse is the matrix

$$A_l^\dagger = (A^T A)^{-1} A^T \quad (\text{A.59})$$

which is termed *left pseudo-inverse*, since $A_l^\dagger A = I_n$. If W_l is an $(m \times m)$ *positive definite* matrix, a *weighted left pseudo-inverse* is given by

$$A_l^\dagger = (A^T W_l^{-1} A)^{-1} A^T W_l^{-1}. \quad (\text{A.60})$$

If A has more columns than rows ($m < n$) and has rank m , a special right inverse is the matrix

$$A_r^\dagger = A^T(AA^T)^{-1} \quad (\text{A.61})$$

which is termed *right pseudo-inverse*, since $AA_r^\dagger = I_m$ ³. If W_r is an $(n \times n)$ *positive definite* matrix, a *weighted* right pseudo-inverse is given by

$$A_r^\dagger = W_r^{-1} A^T (AW_r^{-1} A^T)^{-1}. \quad (\text{A.62})$$

The pseudo-inverse is very useful to invert a linear transformation $y = Ax$ with A a full-rank matrix. If A is a square nonsingular matrix, then obviously $x = A^{-1}y$ and then $A_l^\dagger = A_r^\dagger = A^{-1}$.

If A has more columns than rows ($m < n$) and has rank m , then the solution x for a given y is not unique; it can be shown that the expression

$$x = A^\dagger y + (I - A^\dagger A)k, \quad (\text{A.63})$$

with k an arbitrary $(n \times 1)$ vector and $A^\dagger$ as in (A.61), is a solution to the system of linear equations established by (A.34). The term $A^\dagger y \in \mathcal{N}^\perp(A) \equiv \mathcal{R}(A^T)$ minimizes the norm of the solution $\|x\|$, while the term $(I - A^\dagger A)k$ is the projection of k in $\mathcal{N}(A)$ and is termed *homogeneous solution*.

On the other hand, if A has more rows than columns ($m > n$), the equation in (A.34) has no solution; it can be shown that an *approximate* solution is given by

$$x = A^\dagger y \quad (\text{A.64})$$

where $A^\dagger$ as in (A.59) minimizes $\|y - Ax\|$.

A.8 Singular Value Decomposition

For a nonsquare matrix it is not possible to define eigenvalues. An extension of the eigenvalue concept can be obtained by singular values. Given a matrix A of dimensions $(m \times n)$, the matrix $A^T A$ has n nonnegative eigenvalues $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n \geq 0$ (ordered from the largest to the smallest) which can be expressed in the form

$$\lambda_i = \sigma_i^2 \quad \sigma_i \geq 0.$$

The scalars $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_n \geq 0$ are said to be the *singular values* of matrix A . The *singular value decomposition* (SVD) of matrix A is given by

$$A = U \Sigma V^T \quad (\text{A.65})$$

³ Subscripts l and r are usually omitted whenever the use of a left or right pseudo-inverse is clear from the context.

where U is an $(m \times m)$ orthogonal matrix

$$U = [u_1 \quad u_2 \quad \dots \quad u_m], \quad (\text{A.66})$$

V is an $(n \times n)$ orthogonal matrix

$$V = [v_1 \quad v_2 \quad \dots \quad v_n] \quad (\text{A.67})$$

and Σ is an $(m \times n)$ matrix

$$\Sigma = \begin{bmatrix} D & O \\ O & O \end{bmatrix} \quad D = \text{diag}\{\sigma_1, \sigma_2, \dots, \sigma_r\} \quad (\text{A.68})$$

where $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_r > 0$. The number of nonnull singular values is equal to the rank r of matrix A .

The columns of U are the eigenvectors of the matrix AA^T , whereas the columns of V are the eigenvectors of the matrix $A^T A$. In view of the partitions of U and V in (A.66) and (A.67), it is: $Av_i = \sigma_i u_i$, for $i = 1, \dots, r$ and $Av_i = 0$, for $i = r + 1, \dots, n$.

Singular value decomposition is useful for analysis of the linear transformation $y = Ax$ established in (A.34). According to a geometric interpretation, the matrix A transforms the unit sphere in $\mathbb{R}^n$ defined by $\|x\| = 1$ into the set of vectors $y = Ax$ which define an *ellipsoid* of dimension r in $\mathbb{R}^m$. The singular values are the lengths of the various axes of the ellipsoid. The *condition number* of the matrix

$$\kappa = \frac{\sigma_1}{\sigma_r}$$

is related to the eccentricity of the ellipsoid and provides a measure of ill-conditioning ($\kappa \gg 1$) for numerical solution of the system established by (A.34).

Bibliography

- Boullion T.L., Odell P.L. (1971) *Generalized Inverse Matrices*. Wiley, New York.
- DeRusso P.M., Roy R.J., Close C.M., A.A. Desrochers (1998) *State Variables for Engineers*. 2nd ed., Wiley, New York.
- Gantmacher F.R. (1959) *Theory of Matrices*. Vols. I & II, Chelsea Publishing Company, New York.
- Golub G.H., Van Loan C.F. (1989) *Matrix Computations*. 2nd ed., The Johns Hopkins University Press, Baltimore, Md.
- Noble B. (1977) *Applied Linear Algebra*. 2nd ed., Prentice-Hall, Englewood Cliffs, N.J.