

UNCLASSIFIED
AD 426876

DEFENSE DOCUMENTATION CENTER
FOR
SCIENTIFIC AND TECHNICAL INFORMATION
CAMERON STATION, ALEXANDRIA, VIRGINIA



UNCLASSIFIED

NOTICE: When government or other drawings, specifications or other data are used for any purpose other than in connection with a definitely related government procurement operation, the U. S. Government thereby incurs no responsibility, nor any obligation whatsoever; and the fact that the Government may have formulated, furnished, or in any way supplied the said drawings, specifications, or other data is not to be regarded by implication or otherwise as in any manner licensing the holder or any other person or corporation, or conveying any rights or permission to manufacture, use or sell any patented invention that may in any way be related thereto.

64-6

D1-82-0307

426876

CATALOGED BY DDC

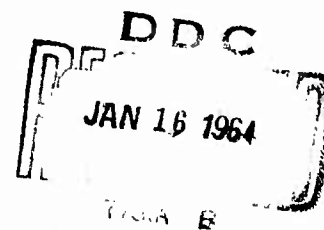
AS AD NO. _____

426876

Also available from the author

BOEING SCIENTIFIC
RESEARCH
LABORATORIES

An Approach to Nonlinear Networks



Bernard Friedman

Mathematics Research

September 1963

AN APPROACH TO NONLINEAR NETWORKS

by

Professor Bernard Friedman
University of California, Berkeley

Mathematical Note No. 325

Mathematics Research Laboratory

BOEING SCIENTIFIC RESEARCH LABORATORIES

September 1963

Chapter I Exterior Algebra

Introduction

This report will bring together several mathematical topics which seem to be relevant to the treatment of nonlinear networks. The main goal will be to discuss some stability theorems recently obtained by Brayton and Moser [1]. Before reaching this goal, we shall present some exterior, or Grassman algebra, some algebraic topology, some theorems about electrical networks, and a few remarks about Liapunov stability. Of course, the final results could be obtained more directly without the excursions into exterior algebra or algebraic topology. However, since these topics do give some additional insight, it was thought worthwhile to include them. Because there are many places where a rigorous and more complete treatment may be found, the approach throughout has been intentionally kept to a very elementary and a very intuitive level.

Linear Spaces

We shall begin with a discussion of the algebra of three-dimensional vectors. Mathematicians have learned that it is more important to know how things behave, than to know what they are. One well-known fact about the behavior of vectors is that they can be added, that is, two vectors a and b can be combined to form a third vector which is called their sum and which is written $a + b$. The operation by which the sum is formed has the following properties:

1. The associative property, that is, $a + (b + c) = (a + b) + c$;

2. a neutral element, written 0 , exists such that for all vectors a , we have $a + 0 = 0 + a = a$;
3. every vector a has an inverse element $-a$ such that $a + (-a) = (-a) + a = 0$;
4. the commutative property, that is, $a + b = b + a$.

These properties are evident from the geometrical definition of addition by the Parallelogram Law which is illustrated in Figure 1.

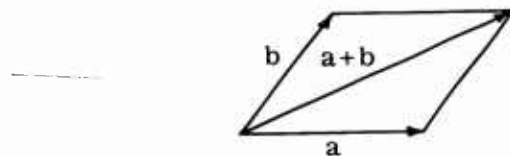


Figure 1

Any set of elements for which an operation such as $+$ can be carried out, and for which the properties 1,2,3, are satisfied is called a group. If Property 4 is also satisfied, the group is said to be commutative.

Besides addition, there is another operation possible with vectors, namely multiplication with a scalar. Let us denote scalars by Greek letters; then the operation of multiplying by a scalar has the following properties:

- | | |
|--|---|
| 5. $\alpha(a + b) = \alpha a + \alpha b$; | 6. $(\alpha + \beta)a = \alpha a + \beta a$; |
| 7. $\alpha(\beta a) = (\alpha\beta)a$; | 8. $1a = a$. |

Any set of elements for which an operation having the properties 1,2,3, and 4 can be defined, and for which an operation having the properties 5,6,7, and 8 can be defined is called a linear space, or a linear vector space. The elements $a_1, a_2, \dots, a_n$ in a linear space are said to be linearly independent if a relation such as

$$\alpha_1 a_1 + \alpha_2 a_2 + \dots + \alpha_n a_n = 0$$

implies that $\alpha_1 = \alpha_2 = \dots = \alpha_n = 0$. If the elements $a_1, a_2, \dots, a_n$ are not linearly independent, they are said to be linearly dependent.

The geometrical significance of linear independence is clear. One vector is linearly independent if, and only if, it is not the zero vector. Two vectors are linearly independent if, and only if, they do not lie in the same line. Three vectors are linearly independent if, and only if, they do not lie in the same plane.

The vectors $a_1, a_2, \dots, a_n$ are said to form a basis of the linear space if every element of the space can be expressed uniquely as a linear combination of the basis, that is, given any element x in the linear space, there exists a unique set of scalars $\alpha_1, \alpha_2, \dots, \alpha_n$ such that

$$x = \alpha_1 a_1 + \alpha_2 a_2 + \dots + \alpha_n a_n.$$

The number of vectors in the basis is called the dimension of the space. The space of geometry and physics obviously has dimension three because every vector can be expressed uniquely as a linear combination of the fundamental $\vec{i}, \vec{j}, \vec{k}$ vectors.

Bivectors

Besides the operations of addition and of multiplying by a scalar, there are other vector operations possible, namely the operations of multiplying two vectors together to obtain the dot product and the cross product. In this section we shall introduce a third type of product, the exterior or wedge product. The wedge product is very closely related to the vector product in a three-dimensional space. However, in higher dimensional spaces, the vector product does not exist but the wedge product always exists. Also, the wedge product has the associative property which the vector product does not.

We shall introduce the wedge product geometrically. A vector has the geometrical significance of a directed line segment which has a magnitude, its length, and a direction. Consider a parallelogram (see Figure 2) having the vectors a and b as sides. This parallelogram has a magnitude,

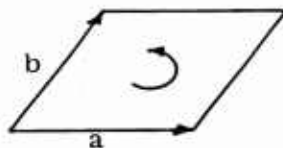


Figure 2

its area, but we shall see that it is useful to assign to the parallelogram also an orientation, that is, a sense or direction in which the sides are traversed. To see the need for an orientation, consider a triangle

(see Figure 3) whose area we shall denote by $\overline{ABC}$. Let O be a point

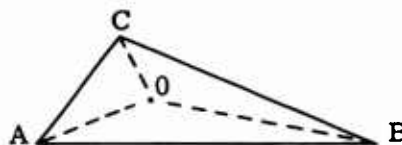


Figure 3

inside the triangle, then

$$\overline{ABC} = \overline{OAB} + \overline{OBC} + \overline{OCA}. \quad (1)$$

However, if the point O is outside the triangle as in Figure 4,

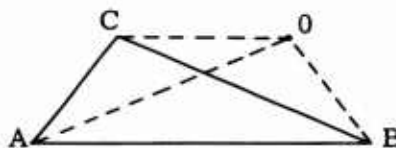


Figure 4

then the equation (1) is no longer true. Instead, we have

$$\overline{ABC} = \overline{OAB} - \overline{OBC} + \overline{OCA}. \quad (2)$$

It would certainly be convenient to be able to use equation (1) no matter where O is placed. This can be done if we agree that an area is to be considered positive or negative according as its vertices are described in a counter-clockwise or a clockwise order. For example, $\overline{ABC}$ will denote the positive number equal to the magnitude of the area of the triangle in Figure 3 or Figure 4, whereas $\overline{ACB}$ will denote the negative number whose absolute value is equal to the magnitude of the area of the triangle. With this interpretation, equation (1) is valid

for both Figures 3 and 4, because in Figure 3 $\overline{OBC}$ describes the triangle in a counter-clockwise sense and thus $\overline{OBC}$ represents a positive number, whereas in Figure 4 $\overline{OBC}$ describes the triangle in a clockwise sense and thus $\overline{OBC}$ represents a negative number. Henceforth, we shall consider all areas to be oriented.

Just as two vectors are considered equal if they have the same length and are parallel in the same direction, so we shall consider two parallelograms equal if they have equal areas, the same orientation, and lie in the same or parallel planes. We shall say these parallelograms represent a bivector. We shall denote by $a \wedge b$ a bivector such as the one represented by the parallelogram in Figure 2, if the orientation of the parallelogram is such that the vector a must be rotated in the counter-clockwise sense to coincide with the vector b . However, if this parallelogram is oriented in the clockwise sense, we denote the bivector by $b \wedge a$. Because the only difference between the bivectors is in their orientation or, what is equivalent, in the sign attached to the area, we shall write

$$a \wedge b = - b \wedge a. \quad (3)$$

Thus, the multiplication indicated in the wedge product is anti-commutative just as is the multiplication in the cross product. Notice that since $a \wedge a = - a \wedge a$, we must have $2 a \wedge a = 0$, or $a \wedge a = 0$. We conclude that the wedge product of a vector with itself is always zero.

It should be emphasized that a wedge product of two vectors is not a vector but it is a bivector which is represented by a parallelogram

with a definite orientation. We can introduce a procedure for adding two bivectors together. This addition procedure is indicated in Figure 5.

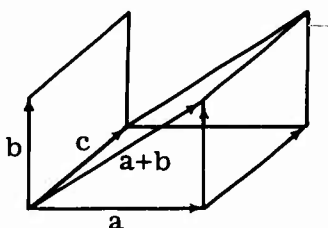


Figure 5

We find that this addition has the properties 1,2,3,4 of the preceding section.

We can also define the operation of multiplying a bivector by a scalar. The product of the bivector $a \wedge b$ by the scalar α , written $\alpha a \wedge b$, is the bivector represented by a parallelogram whose area is $|\alpha|$ times the area of the parallelogram in Figure 2 representing $a \wedge b$, whose orientation is the same or opposite to that in Figure 2 according as α is greater or less than zero, and whose plane is the same or parallel to that of Figure 2. It is easy to see that this operation of multiplying a bivector by a scalar has properties 5,6,7,8 of the preceding section. Consequently, the set of bivectors forms a linear space. We shall soon show that this space is three-dimensional.

Let us illustrate the procedure for finding a wedge product by considering a simple example. Suppose $a = \alpha_1 i + \alpha_2 j + \alpha_3 k$ and $b = \beta_1 i + \beta_2 j + \beta_3 k$, then

$$a \wedge b = \alpha_1 \beta_2 i \wedge j + \alpha_1 \beta_3 i \wedge k + \alpha_2 \beta_1 j \wedge i + \alpha_2 \beta_3 j \wedge k + \alpha_3 \beta_1 k \wedge i + \alpha_3 \beta_2 k \wedge j. \quad (4)$$

Notice we have used the fact that the wedge product of a vector with itself is zero. The result can be still further simplified by using the anti-commutativity (3) of the product. We have $j \wedge i = -i \wedge j$, $k \wedge j = -j \wedge k$, $i \wedge k = -k \wedge i$, and therefore the product in (4) simplifies to

$$a \wedge b = (\alpha_1 \beta_2 - \alpha_2 \beta_1) i \wedge j + (\alpha_2 \beta_3 - \alpha_3 \beta_2) j \wedge k + (\alpha_3 \beta_1 - \alpha_1 \beta_3) k \wedge i. \quad (5)$$

Since every vector can be written as a linear combination of i, j , and k , it is clear that every bivector will be a linear combination of the three bivectors $i \wedge j$, $j \wedge k$, and $k \wedge i$. This shows that the space of bivectors is three-dimensional.

The equation (5) suggests the connection between the wedge product and the vector product. Notice that each of the fundamental bivectors $i \wedge j$, $j \wedge k$, and $k \wedge i$, has a unique normal, namely, the vectors k, i , and j , respectively. If we replace each bivector in (5) by its normal, we get the vector

$$(\alpha_2 \beta_3 - \alpha_3 \beta_2) i + (\alpha_3 \beta_1 - \alpha_1 \beta_3) j + (\alpha_1 \beta_2 - \alpha_2 \beta_1) k,$$

which is the vector product of a and b .

Determinants

The procedure for forming wedge products can be extended to form the wedge product of three vectors. This triple wedge product, which we call a trivector, is represented geometrically by the oriented volume of the parallelepiped generated by the three vectors (see Figure 6). The volume

is oriented by assigning an order in which the vectors $a, b,$ and c are

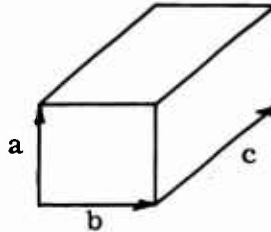


Figure 6

described. If the order is such that the vectors form a right-handed set, the volume is considered positive. If the vectors form a left-handed set, the volume is considered negative.

Let us illustrate the procedure for forming the trivector by considering an example in which the vectors a and b are the vectors used in (4) and the vector $c = \gamma_1 i + \gamma_2 j + \gamma_3 k$. Using (5), we find that

$$a \wedge b \wedge c = \gamma_3 (\alpha_1 \beta_2 - \alpha_2 \beta_1) i \wedge j \wedge k + \gamma_1 (\alpha_2 \beta_3 - \alpha_3 \beta_2) j \wedge k \wedge i + \gamma_2 (\alpha_3 \beta_1 - \alpha_1 \beta_3) k \wedge i \wedge j. \quad (6)$$

Because of (3), we see that $k \wedge i \wedge j = -i \wedge k \wedge j = i \wedge j \wedge k$, and $j \wedge k \wedge i = i \wedge j \wedge k$, therefore the result in (6) can be simplified to the following:

$$a \wedge b \wedge c = [\gamma_1 (\alpha_2 \beta_3 - \alpha_3 \beta_2) + \gamma_2 (\alpha_3 \beta_1 - \alpha_1 \beta_3) + \gamma_3 (\alpha_1 \beta_2 - \alpha_2 \beta_1)] i \wedge j \wedge k.$$

The expression in the bracket is obviously the determinant of the numbers which are the components of the vectors $a, b,$ and c . This result justifies the geometrical interpretation of the trivector as the volume of the parallelepiped in Figure 6.

A wedge product for any two vectors with n components can be defined by a slight extension of the method used to define the wedge product of two vectors in three-dimensional space. For example, if

$$a = (2, 3, -1, 5), \quad b = (0, 1, -2, 4)$$

are two 4-vectors, we first write them as

$$a = 2e_1 + 3e_2 - e_3 + 5e_4, \quad b = e_2 - 2e_3 + 4e_4$$

where e_1, e_2, e_3, e_4 are the basic 4-vectors defined by the equations

$$e_1 = (1, 0, 0, 0), \quad e_2 = (0, 1, 0, 0),$$

$$e_3 = (0, 0, 1, 0), \quad e_4 = (0, 0, 0, 1).$$

We then form the wedge product using the anti-commutativity rule (3) and the fact that the wedge product of two identical vectors is zero. We get

$$a \wedge b = 2e_1 \wedge e_2 - 4e_1 \wedge e_3 + 8e_1 \wedge e_4 - 5e_2 \wedge e_3 + 7e_2 \wedge e_4 + 6e_3 \wedge e_4.$$

The wedge product of any number of n -vectors can be defined by continued multiplication. For example, if

$$c = e_1 - e_2 + 3e_4,$$

$$d = 2e_1 - 2e_3 + e_4,$$

then the wedge product of the vectors a, b , and c is

$$a \wedge b \wedge c = -9e_1 \wedge e_2 \wedge e_3 + 21e_1 \wedge e_2 \wedge e_4 - 6e_1 \wedge e_3 \wedge e_4 - 21e_2 \wedge e_3 \wedge e_4$$

and the wedge product of the vectors a, b, c , and d is

$$a \wedge b \wedge c \wedge d = 75e_1 \wedge e_2 \wedge e_3 \wedge e_4.$$

The number 75 is the value of the determinant of the four vectors $a, b, c,$ and d , that is,

$$\begin{vmatrix} 2 & 0 & 1 & 2 \\ 3 & 1 & -1 & 0 \\ -1 & -2 & 0 & -2 \\ 5 & 4 & 3 & 1 \end{vmatrix} = 75.$$

This connection between the determinant of four 4-vectors and their wedge product suggests the following theorem: An n by n determinant may be evaluated by forming the wedge product of the n column vectors. Of course, the determinant may also be evaluated by forming the wedge product of the row vectors. The proof of these theorems will be found in [2].

The definition of the determinant by means of a wedge product is not very useful for calculating the value of a determinant, but it is useful in obtaining theoretical results about determinants. For example, we shall use this definition to obtain Cramer's rule for the solution of a system of linear equations. Suppose that $a_1, a_2, \dots, a_n, b$ are given n -vectors and that $\xi_1, \xi_2, \dots, \xi_n$ are n unknowns to be found from the system of equations

$$\xi_1 a_1 + \xi_2 a_2 + \dots + \xi_n a_n = b \quad (7)$$

Let us solve for ξ_1 . Since the wedge product of a vector with itself is zero, we may eliminate ξ_2 by multiplying equation (7) with a_2 . In this way we get

$$\xi_1 a_1 \wedge a_2 + \xi_3 a_3 \wedge a_2 + \dots + \xi_n a_n \wedge a_2 = b \wedge a_2.$$

Similarly, multiplying by $a_3, \dots, a_n$ will eliminate the other unknowns and we find

$$\xi_1 a_1 \wedge a_2 \wedge \dots \wedge a_n = b a_2 \wedge \dots \wedge a_n,$$

or

$$\xi_1 = \frac{b a_2 \wedge \dots \wedge a_n}{a_1 \wedge a_2 \wedge \dots \wedge a_n},$$

which is the well-known Cramer's formula.

Derivations

Given an n -dimensional linear space such as the space of n -component vectors, we may form bivectors by taking the wedge product of two vectors, trivectors by taking the wedge product of three vectors, and so on. We shall say that the original elements of the space are elements of degree one, bivectors are elements of degree two, trivectors are of degree three, and so on until we get to elements of degree n . We may form the sum of elements of different degrees in a purely formal fashion thus:

$$3e_1 + e_2 - 2e_1 \wedge e_2 + 6e_1 \wedge e_2 \wedge e_4.$$

Note that this sum cannot be simplified. Since we can also always form the wedge product of any two elements of arbitrary degree, we have an algebra called the exterior or Grassman algebra over the linear space.

An operator D which maps elements of the exterior algebra into elements of the algebra is said to be a linear operator if

$$D(\alpha x + \beta y) = \alpha Dx + \beta Dy,$$

for any elements x and y of the algebra and any scalars α and β . If the operator D also has the property that it maps elements of degree k into elements of degree $k + v$ and if D behaves like a derivative with respect to products, that is, if

$$D(x \wedge y) = (Dx) \wedge y + (-)^{kv} x \wedge (Dy) \quad (8)$$

for x an element of degree k , then D is called a derivation of degree v . For example, in the space of three-dimensional vectors, if the vectors are functions of the time t , then $\frac{d}{dt}$ is a derivation of degree zero. Another example is given by the operator div or $\nabla \cdot$. This operator is a derivation of degree minus one because it maps vectors which are of degree one into scalars which are of degree zero. The product rule (8) is also obeyed as can be seen from the well-known formula

$$\nabla \cdot (a \times b) = (\nabla \times a) \cdot b - a \cdot (\nabla \times b).$$

Notice that we have used here the vector product instead of the wedge product. This is only possible in three-dimensional space.

We shall prove a result which will be useful to us in the next chapter.

Theorem. If D is a derivation of odd degree v , then D^2 is a derivation of degree $2v$.

The proof is obtained by straightforward calculation. We have, by definition,

$$D(x \wedge y) = (Dx) \wedge y + (-)^{kv} x \wedge (Dy)$$

if x is an element of degree k . Then since Dx is an element of degree $k + v$, we get

$$\begin{aligned} D^2(x \wedge y) &= (D^2x) \wedge y + (-)^{(k+v)v} (Dx) \wedge (Dy) + (-)^{kv} [(Dx) \wedge (Dy) + (-)^{kv} x \wedge (D^2y)] \\ &= (D^2x) \wedge y + (-)^{2kv} x \wedge (D^2y), \end{aligned}$$

because $(-)^{kv+v^2} + (-)^{kv} = 0$ since v is odd. Thus D^2 satisfies the product rule (8) and since it maps elements of degree k into elements of degree $k + 2v$ it is a derivation of degree $2v$.

Two final remarks. First, it is sufficient to define a derivation on elements of degree one because by the product rule (8) it can be extended successively to elements of arbitrary degree. Second, no matter what the degree of a derivation D is, if D is applied to a scalar, it gives zero. This follows immediately from (8) if we put $y = \alpha$, a scalar. We find

$$D(\alpha x) = \alpha Dx + (-)^{kv} x \wedge D\alpha,$$

or

$$x \wedge D(\alpha) = 0,$$

for all α . Therefore, $D(\alpha) = 0$.

Chapter II Topology of Networks

Introduction

This chapter will study the geometry of networks with no reference to their electrical properties. We shall assume that the network is a connected set containing b branches and $n + 1$ nodes. We shall assume that each branch connects two nodes and that each node lies on at least two branches (see Figure 1). Further we shall assume that each branch

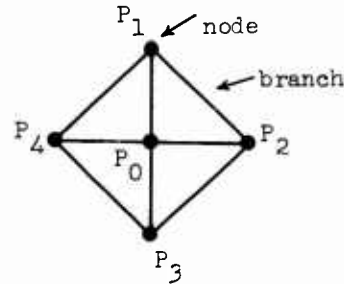


Figure 1

carries a current and a voltage and finally we shall assume that these currents and voltages satisfy Kirchoff's Laws, namely, the sum of the currents at any node point is zero and the sum of the voltages in any closed circuit is also zero. In this chapter we shall assume no relationship between current and voltage. These relationships will be considered in the next chapter where we will discuss the electrical properties of networks.

Chains

Consider a network of b branches and of $n + 1$ nodes. In Figure 1, we have an example of such a network with $b = 8$, and $n = 4$. We shall

name the nodes $P_0, P_1, P_2, \dots, P_n$. The nodes will be called O-cells. A linear combination of O-cells such as

$$5P_0 - 6P_1 + P_2 - P_4$$

will be called a O-chain. Later we shall give a physical interpretation for O-chains. At present, a O-chain will be just a formal sum of the symbols $P_0, P_1, \dots, P_n$. The set of all O-chains forms a linear space.

The branch joining P_i and P_j will be denoted by $P_i P_j$ if the branch is to be considered oriented in the direction from P_i to P_j . Since the direction from P_j to P_i is the opposite of the direction from P_i to P_j , we shall write

$$P_i P_j = -P_j P_i. \quad (1)$$

In this way, an oriented branch may be considered an anti-commutative product of nodes.

A branch will be called a 1-cell and a linear combination of 1-cells such as

$$2P_0 P_1 - 6P_2 P_1 + 5P_4 P_3$$

will be called a 1-chain. The set of all 1-chains forms a linear space. We shall assume that a chain contains only those 1-cells which are branches of the given network. Thus in Figure 1 no chain would contain a 1-cell such as $P_0 P_3$. If we want to emphasize that a 1-chain contains only branches of the network, we shall call it an admissable chain.

A closed circuit on a network will be called a 2-cell. Thus, in Figure 1, $P_0 P_1 P_2$, $P_4 P_3 P_0$, $P_0 P_2 P_3$ are examples of 2-cells. Each 2-cell

will be considered orientated according to the order in which the points are described. Thus $P_0P_1P_2$ is oriented clockwise and $P_0P_2P_1$ is oriented counter-clockwise. We shall write

$$P_0P_1P_2 = - P_0P_2P_1.$$

Notice that this equation is just what would be expected if the 2-cells are considered as triple products of nodes with the products obeying the anti-commutative law (1).

A linear combination of 2-cells such as $2P_0P_1P_2 - P_4P_3P_0 - 3P_0P_2P_3$ will be called a 2-chain. It is convenient to restrict the concept of 2-cells to symbols containing three points only; consequently the circuit $P_1P_2P_3P_4$ in Figure 1 will not be considered as a 2-cell but as the 2-chain $P_1P_2P_3 + P_1P_3P_4$. Because of the orientation of the 2-cells, the sum in the 2-chain contains the common side P_1P_3 traversed twice, once in the positive and once in the negative direction; consequently, the chain reduces to the branches $P_1P_2, P_2P_3, P_3P_4, P_4P_1$ traversed in that order, thus giving the circuit $P_1P_2P_3P_4$.

Boundary

We shall define an algebraic operator on the spaces of chains of all dimensions. This operator will be called the boundary operator, because when this operator is applied to cells it will produce the geometrical boundary of the cell. We denote this operator by the symbol ∂ and we define its action as follows:

The operator ∂ will be a derivation on the space of all chains. When it is applied to a scalar, it gives zero as any derivation will.

When ∂ is applied to a node, it gives one. If we now consider 0-cells as elements of degree one, 1-cells as elements of degree two, and 2-cells as elements of degree three, we may, by the remark at the end of Chapter I, extend ∂ as a derivation from 0-cells first to 1-cells and then to 2-cells. In doing so, we shall use equation (8) of the preceding chapter and we shall assume that ∂ is a derivation of degree minus one.

Let us consider $\partial(P_i P_j)$. We have

$$\partial(P_i P_j) = \partial(P_i)P_j - P_i\partial(P_j) = P_j - P_i. \quad (2)$$

We see that the boundary of the 1-cell $P_i P_j$ is the 0-chain $P_j - P_i$, that is, the end point P_j counted positively and the initial point P_i counted negatively. Note that the boundary of $P_j P_i$ is $P_i - P_j$, which is the negative of the boundary of $P_i P_j$.

Since ∂ is a linear operator, the boundary of a 1-chain may be found by applying ∂ to each term in the chain. Thus

$$\partial(P_1 P_2 + 2P_0 P_1) = P_2 - P_1 + 2P_1 - 2P_0 = P_2 + P_1 - 2P_0.$$

Let us now consider the boundary of the 2-cell $P_i P_j P_k$. Again using (8) of the preceding chapter and (1) and (2) of this chapter, we find that

$$\partial(P_i P_j P_k) = \partial(P_i)P_j P_k - P_i \partial(P_j P_k) = P_j P_k - P_i P_k + P_i P_j.$$

We see that the boundary of the 2-cell $P_i P_j P_k$ in Figure 2 is the 1-chain

$$P_i P_j + P_j P_k - P_i P_k$$

that is, the branches that form the boundary of the 2-cell $P_i P_j P_k$, each

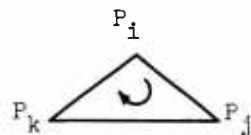


Figure 2

counted with a plus or minus sign according as the orientation of the branch agrees or does not agree with the orientation of the 2-cell.

Again, since ∂ is a linear operator, the boundary of an arbitrary 2-chain may be found by applying ∂ to each term in the chain. Thus the boundary of $P_1 P_2 P_3 + P_1 P_3 P_4$ is

$$P_1 P_2 + P_2 P_3 - P_1 P_3 + P_1 P_3 + P_3 P_4 - P_1 P_4 = P_1 P_2 + P_2 P_3 + P_3 P_4 - P_1 P_4.$$

Notice that this 1-chain actually represents the geometrical boundary of the circuit $P_1 P_2 P_3 P_4$ in Figure 1.

We shall now prove a fundamental topological fact, namely,

Theorem. The boundary of a boundary is zero.

For the proof, we use the theorem we proved at the end of the preceding chapter. Since ∂ is a derivation of odd degree, ∂^2 is a derivation of degree minus two, but $\partial^2(P_i) = \partial(1) = 0$. Therefore, when we use the remark at the end of Chapter I, we conclude that ∂^2 extended to all chains gives zero.

Current Chains

A chain may be interpreted as currents carried in the appropriate cells. For example, a 0-chain such as $\sum \alpha_k P_k$ will describe currents of magnitude α_k flowing into the node P_k . We shall call such a chain a node-current chain. A 1-chain such as $\sum \beta_{ij} P_i P_j$ will describe currents of magnitude β_{ij} flowing along the branch $P_i P_j$. We call such a chain a branch-current chain. A 2-chain such as $\sum \gamma_{ijk} P_i P_j P_k$ will describe currents of magnitude γ_{ijk} flowing in the circuit $P_i P_j P_k$. We call such a chain a loop-current chain.

Consider a branch current $\alpha P_j P_k$. Its boundary is $\alpha P_k - \alpha P_j$ which may be interpreted as a current of magnitude α going into the node P_k and a current of magnitude α going out of the node P_j . Generalizing, we see that the boundary of a branch-current chain is the node-current chain which describes the total amount of current flowing into the nodes because of the branch currents.

We can obtain a similar result for loop-current chains. The boundary of a loop-current chain is the branch-current chain which describes the currents flowing through the branches that form the sides of the loop.

Suppose we have a branch-current chain which is the boundary of a loop-current chain. Using the theorem about the boundary of a boundary being zero, we find that the boundary of this branch-current chain must be a node-current chain which is identically zero. This means that the sum of the currents flowing into any node must be zero; consequently, Kirchoff's Node Law is satisfied. Conversely, if the network is connected,

Kirchoff's Node Law will imply that the branch-current chain is the boundary of a loop-current chain. We shall state this result in the following way: A branch-current chain must be a boundary 1-chain.

Co-chains

To discuss voltages in the same way we have discussed currents, we introduce the concept of a co-chain and a co-boundary. A k -co-chain is a numerical-valued function defined on k -cells. For example, consider the 0-co-chain c_0 defined as follows: The value of c_0 for any $P_j, j \neq 2$, is zero; the value of c_0 for P_2 is one. Note that this co-chain may be extended linearly to be a function of chains; thus,

$$c_0(\sum a_k P_k) = a_2.$$

Similarly, a 1-co-chain is a function defined on 1-cells, and a 2-co-chain is a function defined on 2-cells. Again, these functions may be extended linearly to be functions on chains. We shall say that the value of a p -co-chain c_p on a p -chain d^p is the product of the co-chain with the chain. For example, if the 1-co-chain c_1 has the value a_{ij} on the 1-cell $P_i P_j$, and if the 1-chain d^1 is $\sum \beta_{ij} P_i P_j$ then the product of c_1 and d^1 is

$$c_1 \cdot d^1 = \sum a_{ij} \beta_{ij}. \quad (3)$$

The co-boundary of a $(p-1)$ -co-chain c_{p-1} is the p -co-chain, written δc_{p-1} , which is defined as follows:

$$\delta c_{p-1}(P_{i_0} P_{i_1} \dots P_{i_p}) = c_{p-1}[\partial(P_{i_0} P_{i_1} \dots P_{i_p})]. \quad (4)$$

For example, the co-boundary of the 0-co-chain defined previously is the 1-co-chain whose value on the 1-cell P_{ij} is zero if neither i nor j is two; on the 1-cell $P_{i,2}$ its value is 1; and on the 1-cell $P_{2,j}$ its value is -1. Notice that the co-boundary is zero if P_2 is not on the boundary of the 1-cell, but it is +1 if P_2 is the end point of the 1-cell and -1 if it is the origin of the 1-cell.

We have proved that the boundary of a boundary is zero. Using this result, we can prove the following dual theorem:

The co-boundary of a co-boundary is zero.

The proof is a simple consequence of (4). We have, for any p -chain d^p ,

$$\delta^2 c_{p-2} \cdot d^p = \delta c_{p-2} \cdot \partial d^p = c_{p-2} \cdot \partial^2 d^p = 0$$

because of the fact that $\delta^2 = 0$. Since the value of $\delta^2 c_{p-2}$ on any p -chain is zero, we must have $\delta^2 c_{p-2} = 0$.

Voltage co-chains

Just as currents were defined by chains, so voltages will be defined by co-chains. A 0-co-chain will define the value of a potential at each node; a 1-co-chain will define a voltage in each branch; a 2-co-chain will define a total voltage for a 2-cell. For example, consider the 0-co-chain which states that the potential at P_0 is 2, at P_1 is 4, and at P_2 is 5. We shall write this as

$$2p_0 + 4p_1 + 5p_2.$$

Note that p_k for $k = 0, 1, 2$ is actually the co-chain that has the value 1 for P_k and the value zero for any other 0-cell. Let us consider the co-boundary of this co-chain. It is

$$2p_0p_1 + 3p_0p_2 + p_1p_2.$$

Where $p_i p_j$ is the co-chain that has the value 1 for the 1-cell $P_i P_j$, the value -1 for the 1-cell $P_j P_i$ and the value zero for any other 1-cell. We see that the co-boundary of a node-potential co-chain is the branch-voltage co-chain such that the voltage in each branch is the difference of the node potentials at the ends of the branch.

A similar calculation will show that the co-boundary of a branch-voltage co-chain is a loop-voltage co-chain such that the voltage in each loop is the oriented sum of the voltages in the branches bounding the loop.

Suppose we have a branch-voltage co-chain which is the co-boundary of a node-potential co-chain. Since the co-boundary of a co-boundary is zero, we see that the co-boundary of the branch-voltage co-chain must be zero. This result implies that the sum of the voltages in any loop must be zero. Thus, Kirchhoff's Circuit Law is satisfied. Conversely, it can be shown that if the network is connected and if Kirchhoff's Law is satisfied, then the branch-voltage co-chain must be the co-boundary of a node-potential co-chain. We shall state this result as follows:
The branch voltages must form a co-boundary co-chain.

Tellegen's Theorem

We shall prove an important theorem connecting the branch currents and the branch voltages in an arbitrary network if Kirchhoff's Laws are

satisfied. The theorem, first proved by Tellegen, is as follows:

Let I be a current vector, that is, let I be a b -component vector such that each component represents the current in a particular branch of the network. Let V be a voltage vector, that is, let V be a b -component vector such that each component represents the voltage in a branch. If Kirchhoff's Laws are satisfied, then the scalar product of I and V must be zero.

Notice that no relation between I and V is assumed. All that is required is that the current and voltage vector each separately satisfy Kirchhoff's Laws. The proof of the theorem will follow easily from the topological characterization of branch-currents and branch-voltages given previously. To the vector V there corresponds a co-boundary 1-co-chain c_1 and to the vector I there corresponds a boundary 1-chain d^1 . From (3) we see that the scalar product of I and V is equal to $c_1 \cdot d^1$. Since c_1 is a co-boundary, it may be written as δc_0 and since d^1 is a boundary, it may be written as ∂d^2 . Now by (4),

$$\delta c_0 \cdot \partial d^2 = c_0 \cdot \partial^2 d^2 = 0,$$

because the boundary of a boundary is zero; therefore the theorem is proved.

Tellegen's theorem can be generalized and written in an integral form. Suppose we consider the behavior of a network as a function of some parameter such as the time t . For each value of t in some range T there will be a current vector $I(t)$ and a voltage vector $V(t)$. We shall prove the following:

Integral Theorem

If I and V independently satisfy Kirchhoff's Laws, then

$$\int_{\Gamma} V dI = 0.$$

The proof is an easy consequence of Tellegen's theorem. Since $I(t + h)$ and $I(t)$ are both current vectors, we have

$$V(t)[I(t + h) - I(t)] = 0$$

for all values of t and h . Letting h go to zero, we find that

$$V(t)dI(t) = 0$$

for all values of t . Approximating the integral in the theorem by a sum, we finally obtain the result

$$\int_{\Gamma} V dI = 0.$$

Notice that a similar argument will prove the dual result, namely,

$$\int_{\Gamma} IdV = 0.$$

Chapter III Non-linear Networks

Relations between currents and voltages

In the preceding chapter we have considered the branch currents and the branch voltages as completely independent entities, limited only by Kirchhoff's Laws. However, in any network the current in a branch is related to the voltage in that branch by an electrical device such as a resistor, an inductor, or a capacitor. Kirchhoff's Laws together with the electrical relations between the branch currents and the branch voltages will give us a complete set of equations from which the electrical behavior of the network can be determined.

We begin by specifying the relation between the orientation of the branch and the signs of the branch current and the branch voltage. If the current through a branch is in the same direction as the orientation, the current will be called positive; otherwise, the current will be called negative. The voltage across a branch will be equal to the potential at the endpoint of the branch minus the potential at the initial point of the branch. Since the current goes from a node of higher potential to one of lower potential, our conventions about signs will produce an extra minus sign in the customary relations between current and voltage. For example, consider Figure 1 in which the upper node of the branch is at six volts potential above the potential of the lower node and in which the resistance is two ohms. Of course, the current through the branch is three amperes flowing downward. If the branch is oriented downward, then the current is positive and the voltage is negative. However, if the branch is oriented upward, then

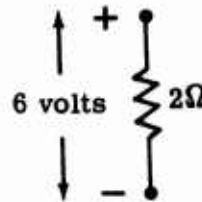


Figure 1

the current is negative and the voltage is positive. In both cases, the relation between the voltage v across the branch and the current i through it is $v = - Ri$, where R is the resistance of the branch.

In addition to linear resistors, we shall also consider non-linear resistors in which the relation between current and voltage is given by the equation $f(i,v) = 0$. This equation can be solved for either of the variables i or v as a function of the other variable. We shall not require that either of the resulting functions be single-valued because we shall permit the characteristic of the resistor to have the shape shown in Figure 2.

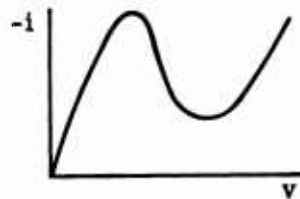


Figure 2

The inductors and capacitors may also be non-linear. For an inductor, the relation between current and voltage may have the form

$$v = - L(i) \frac{di}{dt}$$

where $L(i)$ is a non-negative inductance. For a capacitor, the relation between current and voltage may have the form

$$i = - C(v) \frac{dv}{dt}$$

where $C(v)$ is a non-negative capacitance.

Equilibrium equations

Let us consider again a connected network containing $n + 1$ nodes and b branches. A set of branches which is connected and contains no closed loop is called a tree. A maximal tree is a tree which is not contained in any other tree. It is easy to see that a maximal tree must contain all the nodes; for, if a node did not belong to the tree, a branch joining this node to the tree could be added without spoiling the tree shape. It is also easy to see that a maximal tree must have exactly n branches; for, starting with one fixed node, we must add one branch for each node that is to be connected to this fixed node.

The network will have many maximal trees. Choose one such tree and consider the $b - n = L$ branches that do not belong to the tree. These L branches are called links of the tree. Note that if any link is added to a maximal tree the resulting network will have a loop in it.

From the definition of a tree it is clear physically that the

voltages in the n branches of a maximal tree can be chosen arbitrarily because the tree has no loops. Then by adding each link in turn to the tree and using Kirchoff's Circuit Law we may find the voltage in that link. Since the network contains L links, this argument shows that there are L independent loop equations.

Instead of starting with voltages, we may begin with currents. As we add a link to the maximal tree, we may impress an arbitrary current on it. It is clear that when all the links with their currents have been added to the tree, the currents in the other n branches of the tree can be determined. This argument implies that there are exactly n independent mesh equations. Notice that so far we have not used the electrical properties of the branches but have used only Kirchoff's Laws.

Suppose we pick a set of $r \leq L$ branches whose currents we shall call $\tilde{i}_1, \dots, \tilde{i}_r$ and a set of $s \leq n$ branches whose voltages we shall call $\tilde{v}_{r+1}, \dots, \tilde{v}_{r+s}$. We assume that no branch has both a current and a voltage specified. We shall show that all the other currents and voltages can be determined from these specified currents and voltages. First, from Kirchoff's Laws we have L equations for the voltages and n equations for the currents. Then, in the $b-r-s$ branches for which neither current nor voltage was specified, we use the electrical properties of these branches to obtain $b-r-s$ equations. Thus, we have a total of $L + n + b - r - s = 2b-r-s$ equations to find the unspecified $b-r$ currents and the $b-s$ voltages. Since the number of equations is equal to the number of unspecified quantities, we assume that they can be solved to

give equations of the form

$$\begin{aligned} i_{\mu} &= F_{\mu}(\tilde{i}, \tilde{v}), & 1 \leq \mu \leq b. \\ v_{\mu} &= G_{\mu}(\tilde{i}, \tilde{v}). \end{aligned} \quad (1)$$

Notice that in this discussion we have not used the electrical properties of the branches in which either the current or the voltage was specified. If we assume that the voltage in the r links is related to the current by equations of the form

$$v_{\rho} = g_{\rho}(i_{\rho}), \quad 1 \leq \rho \leq r$$

and if we assume that the current in the s branches is related to the voltage by equations of the form

$$i_{\sigma} = f_{\sigma}(v_{\sigma}), \quad r+1 \leq \sigma \leq r+s$$

we obtain the following $r+s$ equations from which the equilibrium state of the network can be determined:

$$\begin{aligned} g_{\rho}(\tilde{i}_{\rho}) &= G_{\rho}(\tilde{i}, \tilde{v}), & 1 \leq \rho \leq r, \\ f_{\sigma}(\tilde{v}_{\sigma}) &= F_{\sigma}(\tilde{i}, \tilde{v}), & r+1 \leq \sigma \leq r+s. \end{aligned} \quad (2)$$

An example may clarify the argument. Consider the network in Figure 3 for which $n = 2$ and $b = 4$. Branches numbered 3 and 4 will form a maximal tree and branches 1 and 2 will be links for this tree. Let us assume that the currents, $\tilde{i}_1$ and $\tilde{i}_2$ in branches 1 and 2 are specified and also the voltage $\tilde{v}_3$ in branch 3. This implies that

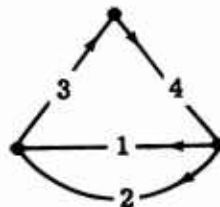


Figure 3

$r = 2$ and $s = 1$. By the use of Kirchoff's Laws we find that the currents in branches 3 and 4 are given by the equations $i_3 = \tilde{i}_1 + \tilde{i}_2$, $i_4 = \tilde{i}_1 + \tilde{i}_2$. Next, assuming that the electrical properties of the unspecified branch 4 are such that the relation between current and voltage is given by $v = g_4(i)$ we find that

$$\begin{aligned} v_4 &= g_4(i_4) = g_4(\tilde{i}_1 + \tilde{i}_2) \\ \text{and} \\ v_1 &= v_2 = -v_4 - \tilde{v}_3 = -\tilde{v}_3 - g_4(\tilde{i}_1 + \tilde{i}_2). \end{aligned}$$

Thus, all the unspecified variables have been expressed in terms of the specified ones. Now, using the electrical properties of the specified branches, we finally obtain the following equilibrium equations for the network:

$$\begin{aligned} g_1(\tilde{i}_1) &= g_2(\tilde{i}_2) = -v_3 - g_4(\tilde{i}_1 + \tilde{i}_2) \\ f_3(\tilde{v}_3) &= +\tilde{i}_1 + \tilde{i}_2. \end{aligned}$$

In writing this set of equilibrium equations we have assumed that all the branches were resistors. Suppose, however, that branches 1 and 2 were inductances and that branch 3 was a capacitor. Then the dynamical equations for the network would be

$$\begin{aligned} L_1(\tilde{i}_1) \frac{d\tilde{i}_1}{dt} &= L_2(\tilde{i}_2) \frac{d\tilde{i}_2}{dt} = +\tilde{v}_3 + g_4(\tilde{i}_1 + \tilde{i}_2), \\ -C_3(\tilde{v}_3) \frac{d\tilde{v}_3}{dt} &= \tilde{i}_1 + \tilde{i}_2. \end{aligned} \tag{3}$$

Notice that any set of currents and voltages can be specified, as long as Kirchoff's Laws are not violated. The only reason for specifying a particular set of currents and voltages is that we wish to write the equilibrium equations of the network in terms of this particular set.

Mixed Potential

A careful study of the dynamical equations for a network shows that they can be written in a simple form. For example, equations (3) may be written as follows:

$$L_1(\tilde{i}_1) \frac{d\tilde{i}_1}{dt} = \frac{\partial P}{\partial \tilde{i}_1}, \quad L_2(\tilde{i}_2) \frac{d\tilde{i}_2}{dt} = \frac{\partial P}{\partial \tilde{i}_2}, \quad -C_3(v_3) \frac{d\tilde{v}_3}{dt} = \frac{\partial P}{\partial \tilde{v}_3},$$

where

$$P = \tilde{v}_3(\tilde{i}_1 + \tilde{i}_2) + \int_0^{(\tilde{i}_1 + \tilde{i}_2)} g_4(i) di.$$

We shall call the scalar function P the mixed potential of the network.

Consider again a network containing $n + 1$ nodes and b branches. Suppose that the branches numbered $1, \dots, r$ are inductors and the branches numbered $r + 1, \dots, s$ are capacitors. Just as in the preceding section, we can specify the currents in the inductive branches and the voltages in the capacitive branches. Again, we obtain the equations (1), namely,

$$\begin{aligned} i_\mu &= F_\mu(\tilde{i}, \tilde{v}) \\ v_\mu &= G_\mu(\tilde{i}, \tilde{v}) \end{aligned} \quad 1 \leq \mu \leq b. \quad (4)$$

The currents and voltages are functions of the time t . Let Γ be a curve in the space of currents and voltages representing the variation of these quantities. By the Integral Theorem of Chapter II, we have

$$\begin{aligned} 0 &= \sum_{\mu=1}^b \int_{\Gamma} v_\mu di_\mu = \left(\sum_{\mu=1}^r + \sum_{\mu=r+1}^{r+s} + \sum_{\mu=r+s+1}^b \right) \int_{\Gamma} v_\mu di_\mu \\ &= \sum_{\mu=1}^r \int_{\Gamma} v_\mu di_\mu - \sum_{\mu=r+1}^{r+s} \int_{\Gamma} i_\mu dv_\mu + \left[\sum_{\mu=r+1}^{r+s} i_\mu v_\mu \Big|_{\Gamma} + \sum_{\mu=r+s+1}^b \int_{\Gamma} v_\mu di_\mu \right]. \end{aligned} \quad (5)$$

Consider the two terms in the bracket on the right-hand side. Using (4) for $r + 1 \leq \mu \leq s$, we can express the sum in the bracket as a function of the specified variables $\tilde{i}, \tilde{v}$ at the end of Γ and at the beginning of Γ . For the integral term in the bracket, notice that the branches involved are only the resistive ones in which v_μ depends only on i_μ . (It is immaterial that the relation between i_μ and v_μ may not be single-valued. The initial point of Γ determines a point on the curve $f_\mu(i_\mu, v_\mu) = 0$ and we just integrate along that curve.) The point of the argument is that a term such as

$$\int_{\Gamma} v_\mu di_\mu = \int_{\Gamma} g_\mu(i_\mu) di_\mu$$

depends only on the endpoint of Γ and by (4) this is a function of the specified variables only. Our conclusion is that the bracket term in (5) is a function $P(\tilde{i}, \tilde{v})$ of the specified variables only. We call this function the mixed potential of the network.

We may then write (5) as follows:

$$\sum_{l=1}^r \int_{\Gamma} v_\mu di_\mu - \sum_{r+1}^{r+s} \int_{\Gamma} i_\mu dv_\mu + P(\tilde{i}, \tilde{v}) = 0. \quad (6)$$

Since the variables i, v are independent, we see that (6) implies that

$$\begin{aligned} \tilde{v}_\mu &= - \frac{\partial P}{\partial \tilde{i}_\mu}, & 1 \leq \mu \leq r \\ \tilde{i}_\mu &= \frac{\partial P}{\partial \tilde{v}_\mu}, & r + 1 \leq \mu \leq r + s. \end{aligned}$$

But $\tilde{v}_\mu$, $1 \leq \mu \leq r$, is the voltage across an inductor and $\tilde{i}_\mu$, $r + 1 \leq \mu \leq s$ is the current across a capacitor. Using the

electrical relations for these branches, we finally obtain the dynamical equations of the network, namely

$$\begin{aligned} L_{\mu}(i_{\mu}) \frac{di_{\mu}}{dt} &= \frac{\partial P}{\partial i_{\mu}}, & 1 \leq \mu \leq r \\ -C_{\mu}(v_{\mu}) \frac{dv_{\mu}}{dt} &= \frac{\partial P}{\partial v_{\mu}}, & r+1 \leq \mu \leq r+s, \end{aligned} \quad (7)$$

where

$$P(i,v) = \sum_{r+s+1}^b \int_{\Gamma} v_{\mu} di_{\mu} + \sum_{r+1}^{r+s} i_{\mu} v_{\mu} \Big|_{\Gamma}. \quad (8)$$

For convenience, we have dropped the tilde signs.

The Mixed Potential for a Complete Set

The mixed potential can be given a particularly useful form in case the r currents through the inductors and the s voltages through the capacitors form a complete set of variables. This means they are a set of currents and voltages which do not violate Kirchhoff's Laws and which are such that either the current or the voltage, or both, in every branch is determined by Kirchhoff's Laws only. For example, i_1, i_2 , and v_3 form a complete set of variables for the network in Figure 3 because, no matter what values they have, they do not violate Kirchhoff's Laws because the current in the unspecified branch 4 is determined by Kirchhoff's Laws to be $\tilde{i}_1 + \tilde{i}_2$.

Consider a network in which the specified variables form a complete set. Let N_v be the set of branches of the network in which the voltages

are specified by Kirchoff's Laws and let N_i be the complementary set of branches, that is, those in which the current is specified by Kirchoff's Laws. For example, in Figure 3, N_v would contain only the branch 3 and N_i would contain branches 1,2, and 4. Using this decomposition of the network, we may write (8) as follows:

$$P(i,v) = \left[\sum_{\substack{\mu > r+s \\ \mu \in N_i}} + \sum_{\substack{\mu > r+s \\ \mu \in N_v}} \right] \int_{\Gamma} v_{\mu} di_{\mu} + \sum_{r+1}^{r+s} i_{\mu} v_{\mu} \Big|_{\Gamma}.$$

The second sum in the bracket can be integrated by parts to give

$$P(i,v) = \sum_{\substack{\mu > r+s \\ \mu \in N_i}} \int_{\Gamma} v_{\mu} di_{\mu} - \sum_{\substack{\mu > r+s \\ \mu \in N_v}} \int_{\Gamma} i_{\mu} dv_{\mu} + \sum_{\mu \in N_v} i_{\mu} v_{\mu} \Big|_{\Gamma}. \quad (9)$$

Notice we have used the fact that the inductive branches belong to N_i and the capacitative branches belong to N_v to simplify one of the sums in (9).

The individual terms in (9) can be given an interesting physical interpretation. The first term

$$\int_{\Gamma} v_{\mu} di_{\mu}$$

is well-defined as a line integral over the curve connecting i and v in the branch μ . We call this integral the current potential of the branch μ . Similarly, the integral

$$\int_{\Gamma} i_{\mu} dv_{\mu}$$

is called the voltage potential of the branch μ . Note that in (9) the current potential is evaluated only for branches in N_i and the voltage

potential only for branches in N_v . Since the current is specified in N_i and the voltage in N_v , we see that

$$P(i,v) = F(i) - G(v) + \sum_{\mu \in N_v} i_{\mu} v_{\mu} \Big|_{\Gamma}, \quad (10)$$

where $F(i)$ is the current potential and $G(v)$ is the voltage potential. It can be shown that the third term on the right-hand side of (10) is a bilinear function of the specified currents and voltages. For details we refer to [1].

Example of the Calculation of the Mixed Potential

Consider again the network illustrated in Figure 3, but now we shall specify its electrical properties as illustrated in Figure 4. The symbol

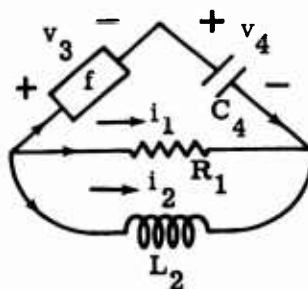


Figure 4

in branch 3 indicates the presence of a non-linear resistor whose characteristic will be assumed to be of the form $v_3 = -f(i_3)$.

Since there is only one inductor and one capacitor in the network of Figure 4, the appropriate variables to specify the dynamical behavior of this system are i_2 , the current across the inductor, and v_4 , the

voltage across the capacitor. Note that i_2 and v_4 do not form a complete set of variables because the use of Kirchhoff's Laws alone without the electrical properties of the network does not determine the value of either i_3 or v_3 .

Nevertheless, even though the set of variables is not complete, a mixed potential does exist and the dynamical equations for the network in Figure 4 can be written in the form (7). To see this, let us write Kirchhoff's Laws using the electrical properties of the network. It is clear that $i_3 = -i_1 - i_2$. We have

$$\begin{aligned} L_2 \frac{di_2}{dt} &= i_1 R_1 \\ v_4 &= f(i_1 + i_2) + i_1 R_1 \\ C_4 \frac{dv_4}{dt} &= -(i_1 + i_2). \end{aligned} \tag{11}$$

Since we have three equations for the three variables i_1, i_2, v_4 , the system is completely determined.

Notice that the second equation of the set (11) is not a differential equation. This occurrence is a result of the fact that the set of variables i_2 and v_4 is not a complete set. We can of course eliminate i_1 from the equations (11) and thus obtain a set of differential equations only. To do this, we notice that the second equation of (11) may be written

$$v_4 + i_2 R_1 = f(i_1 + i_2) + (i_1 + i_2) R_1, \tag{12}$$

so that $v_4 + i_2 R_1$ is a function of $i_1 + i_2$. Let us invert this functional relationship to get $i_1 + i_2$ as a function of $v_4 + i_2 R_1$ and let us write this new relation as

$$i_1 + i_2 = g(v_4 + i_2 R_1).$$

Now eliminating i_1 from the first and third equations of (11), we get the desired result:

$$\begin{aligned} C_4 \frac{dv_4}{dt} &= -g(v_4 + i_2 R_1) \\ L_2 \frac{di_2}{dt} &= R_1 g(v_4 + i_2 R_1) - R_1 i_2. \end{aligned} \tag{13}$$

The set of equations (13) may be written in the form (7) if we introduce the mixed potential P defined as follows:

$$P = \int_0^{(v_4 + i_2 R_1)} g(v) dv - \frac{1}{2} R_1 i_2^2. \tag{14}$$

Notice that the form (14) for the mixed potential is not in the form (10), that is, it is not the sum of a term depending only on the voltage plus a term depending only on the current plus a term which is bilinear in the current and the voltage. Instead, the first term on the right-hand side of (14) is a complicated function of both the current and the voltage.

If we modify the network in Figure 4 so that it is the network illustrated in Figure 5, we shall have an example of a network containing a complete set of variables and we shall find a mixed potential of the form (10). Since there are two inductors and one capacitor in Figure 5,

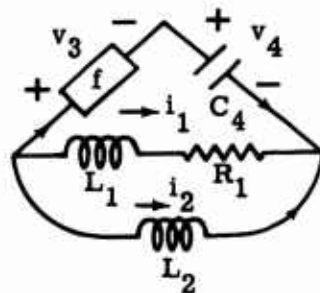


Figure 5

the appropriate variables to specify the dynamical behavior are i_1, i_2 , and v_4 , and these variables form a complete set. Again, $i_3 = -i_1 - i_2$, and Kirchoff's Laws for this network become

$$\begin{aligned}
 L_2 \frac{di_2}{dt} &= L_1 \frac{di_1}{dt} + i_1 R_1 \\
 L_2 \frac{di_2}{dt} &= v_4 - f(i_1 + i_2) \\
 C_4 \frac{dv_4}{dt} &= -(i_1 + i_2) \\
 L_1 \frac{di_1}{dt} &= v_4 - f(i_1 + i_2) - i_1 R_1.
 \end{aligned} \tag{15}$$

Notice that the fourth equation of this set is implied by the first and second equations. Considering the last three equations of this set, we have a set of independent equations for the three variables i_1, i_2 , and v_4 . These equations are all differential equations which can be written in the form (7) if we take

$$P = -\frac{1}{2} R_1 i_1^2 - \int_0^{(i_1+i_2)} f(i) di + v_4 (i_1 + i_2). \tag{16}$$

Equation (16) shows that P has the form (10) with the current potential equal to

$$-\frac{1}{2} R_1 i_1^2 - \int_0^{(i_1+i_2)} f(i) di,$$

no voltage potential term, and the remaining part being the bilinear function $v_4(i_1 + i_2)$. The formula (16) can be written down immediately from (10) without going through (15). From (9) and (10), $F(i)$ is the current potential evaluated along the resistors in the branches in which the current is defined by Kirchoff's Laws alone without the electrical properties of the network. In Figure 5, the current is defined in all the branches but since there are resistors only in branches 1 and 3, we find

$$F(i) = -\frac{1}{2} R_1 i_1^2 - \int_0^{(i_1+i_2)} f(i) di. \quad (17)$$

The voltage potential $G(v)$ must be evaluated along the resistors in the branches in which the voltage is defined by Kirchoff's Laws alone. The only such branch is 4 and since this branch has no resistance, we have $G(v) = 0$.

The remaining part of the mixed potential is defined in (9) as

$$\sum_{\mu \in N_v} i_\mu v_\mu \Big|_\Gamma, \quad (18)$$

where the sum is to be taken over all branches in which Kirchoff's Laws specify the voltage. Again, the only such branch is 4, and since the current in this branch is $-(i_1 + i_2)$ and the voltage is $-v_4$ because of the assigned orientation of the branch, we find that (18) equals

$v_4(i_1 + i_2)$. Combining (17) and (18), we get

$$P = -\frac{1}{2} R_1 i_1^2 - \int_0^{(i_1+i_2)} f(i) di + v_4(i_1 + i_2),$$

in agreement with (16).

Vector Form for the Dynamical Equations

Consider a network with a complete set of variables $i_k, \dots, i_r, v_{r+1}, \dots, v_{r+s}$. Denote this set of variables by the $(r + s)$ -component vector x . Let L denote the diagonal matrix whose entries are the inductances $L_1(i), \dots, L_r(i)$ and let C denote the diagonal matrix whose entries are $C_{r+1}(v), \dots, C_{r+s}(v)$. Put

$$J = \begin{pmatrix} -L & 0 \\ 0 & C \end{pmatrix},$$

where the zeros denote rectangular matrices whose entries are all zero.

If a dot denotes differentiation with respect to time, we may write (7)

in the following vector form:

$$-J \dot{x} = \frac{\partial P(x)}{\partial x}, \quad (19)$$

where $\frac{\partial P}{\partial x}$ denotes the gradient of the function P .

Consider any solution of the equations (19). This solution is defined by specifying x as a function of t ; say, $x = X(t)$. Note this specification defines a curve Γ . Let us consider the way in which the mixed potential varies along Γ . We have

$$\frac{dP(X(t))}{dt} = \frac{\partial P}{\partial x} \cdot \frac{dx}{dt} = -\langle J\dot{X}, \dot{X} \rangle. \quad (20)$$

The expression on the right-hand side of (20) is a quadratic form in the vector $\dot{X}(t)$. Suppose that J is a positive-definite matrix, then (20) shows that the value of the mixed potential always decreases along any curve Γ . Considering the mixed potential as some kind of "energy" of the network, we are led to believe that, because the "energy" always decreases, the network must always be stable. This semi-intuitive reasoning will be made precise in the next section where we discuss Liapunov's Stability Theorem.

Liapunov Stability

Let x denote an n -component vector and let $F(x)$ be a continuous and differentiable vector-valued function of x . Consider the vector system of first-order equations

$$\dot{X} = F(X). \quad (21)$$

The equilibrium points of this system are the points x in n -space E_n at which $F(x) = 0$. An invariant set of (21) is a set of points in E_n such that if a solution of (21) starts in the invariant set at $t = 0$, then for all future time the solution remains in the invariant set. It is clear that the equilibrium points of (21) form an invariant set.

Let $V(\vec{x})$ be a scalar-valued function of x with continuous first partial derivatives. Along any solution of (21) specified by $x = X(t)$, we have

$$\frac{dV(X(t))}{dt} = F(X(t)) \cdot \text{grad } V.$$

Let us put

$$\dot{V}(x) = F(x) \cdot \text{grad } V(x),$$

so that

$$\frac{dV(X(t))}{dt} = \dot{V}(X(t)). \quad (22)$$

If $\dot{V}(x) \leq 0$ and $V(x) \geq 0$ for some region of E_n containing the origin, then $V(x)$ will be called a Liapunov function. We shall prove the following [3]:

Theorem . . Let R_C denote the region around the origin in E_n where $V(x) < C$. Suppose that R_C is bounded and that inside R_C $V(x) > 0$ for $x \neq 0$, and $\dot{V}(x) \leq 0$. Let M be the largest invariant set of (21) contained in the set of points for which $\dot{V}(x) = 0$; then, every solution of (21) which starts in R_C at $t = 0$ tends to M as t goes to infinity.

To prove the theorem, we consider any solution $X(t)$ of (21) which is in R_C at $t = 0$; therefore $V(X(0)) < C$. Because of (22) and the conditions on $V(x)$, the value of V along this solution is non-increasing for all $t \geq 0$, and $V(X(t)) < C$ for all $t \geq 0$; we conclude that the solution $X(t)$ can never reach the boundary of R_C at which $V(x) = C$ and therefore $X(t)$ must remain in R_C . Since $V(X(t))$ is non-increasing and since $V(x) \geq 0$ for all points in R_C , we see that $V(X(t))$ must converge to a limit, say C_0 , as t goes to infinity.

Let L_+ be the limiting set of $X(t)$, that is, the set of points p in E_n for which there exist an infinite sequence of positive times

$t_1, t_2, \dots$, such that $X(t_k)$ converges to p as k goes to infinity. Note that the limiting set must be an invariant set of (21). For, suppose L_+ were not an invariant set. Then there would exist a solution $\tilde{X}(t)$ of (21) starting for $t = 0$ at a point p_0 in L_+ and such that $\tilde{X}(t)$ is not in L_+ for $t > 0$ and t small enough. Since $F(x)$ is continuously differentiable, it can be approximated for x in the neighborhood of p_0 by $F(p_0) + K \cdot (x - p_0)$ where K is a matrix with constant entries; consequently, for x near enough to p_0 , the solutions of

$$\dot{X} = F(X)$$

will be approximated by the solutions of

$$\dot{X} = F(p_0) + K \cdot (X - p_0). \quad (23)$$

However, for (23), it is easy to show that a solution such as $\tilde{X}(t)$ cannot exist and then, by continuity, it follows that $\tilde{X}(t)$ does not exist for (21). Therefore, L_+ is an invariant set.

From the preceding discussion about the limit of $V(X(t))$ as t goes to infinity, it follows that $V(x) = C_0$ for all x in L_+ ; consequently, $\dot{V} = 0$ on L_+ . Also, the set L_+ is in R_C . Therefore, if M is the largest invariant set contained in the set of points at which $\dot{V} = 0$, then M contains L_+ , the limiting set of the solution $X(t)$. This result proves the theorem.

It is useful to make a slight generalization of Theorem I as follows:

Theorem II. Suppose for all x in E_n that $V(x) > 0$ for $x \neq 0$ and that $\dot{V}(x) \leq 0$; suppose further that $V(x) \rightarrow \infty$ as $x \rightarrow \infty$; then every solution of (21) tends to the set M of Theorem I.

The proof follows easily from Theorem I. Consider any solution $X(t)$ of (21) and suppose that $V(X(0)) = C'$. By the hypothesis of the theorem the set of points in E_n such that $V(x) < C = C' + \epsilon$, for arbitrary $\epsilon > 0$, is bounded. Now, we can use Theorem I to complete the proof.

Application to Networks

From equation (20), we see that $P(x)$ will be a Liapunov function if J is positive definite and then Theorem II could be applied to give a stability theorem. In general, J is not positive definite and other arguments must be used. We shall conclude with a statement of a stability theorem for some networks [1].

Notice that if J_1 is a matrix and $P_1(x)$ a scalar-valued function such that the solutions of (19) satisfy

$$-J_1 \dot{X} = \frac{\partial P_1(x)}{\partial x}, \quad (24)$$

then, because of linearity, the solutions of (19) satisfy all equations of the form

$$-\tilde{J} \dot{X} = \frac{\partial \tilde{P}(x)}{\partial x},$$

where

$$\begin{aligned} \tilde{J} &= \alpha J + \beta J_1 \\ \tilde{P} &= \alpha P + \beta P_1 \end{aligned} \quad (25)$$

with α and β arbitrary constants.

A possible choice for J_1, P_1 in (24) is the following:

$$J_1 = P_{xx} S J, P_1 = \frac{1}{2} \text{grad } P \cdot S \text{ grad } P, \quad (26)$$

where S is a symmetric matrix with arbitrary constants for entries.

To see this, start from (19), namely,

$$-J \dot{X} = \text{grad } P,$$

and multiply in front by the matrix $P_{xx} S$. We get

$$-J_1 \dot{X} = -P_{xx} S J \dot{X} = P_{xx} \cdot S \text{ grad } P$$

but notice that

$$\begin{aligned} \frac{1}{2} \frac{\partial}{\partial x} (\text{grad } P \cdot S \text{ grad } P) &= \frac{1}{2} P_{xx} \cdot S \text{ grad } P + \frac{1}{2} \text{grad } P \cdot S P_{xx} \\ &= P_{xx} \cdot S \text{ grad } P, \end{aligned}$$

because S is symmetric. This justifies (26).

We now state a theorem about the stability of networks. For the proof, we refer to [1, page 38].

Theorem III. Suppose

$$P(x) = -\frac{1}{2}(i, Ai) + B(v) + i \cdot (\gamma v - a),$$

where A is a constant symmetric matrix, γ is an arbitrary constant matrix, and a is a constant vector. Suppose A is positive definite,

$B(v) + |\gamma v| \rightarrow \infty$ as $|v| \rightarrow \infty$, and the norm of the matrix

$$L^{\frac{1}{2}}(i)A^{-1}\gamma C^{-\frac{1}{2}}(v)$$

is less than $1 - \delta$, with $\delta > 0$, for all i, v ; then, as $t \rightarrow \infty$, all the solutions of (19) tend to the equilibrium points, that is, the points at which $\text{grad } P = 0$.

The idea of the proof is to find a suitable matrix S in (26) and suitable α and β in (25) so that the conditions of Theorem II are satisfied. The details are in [1].

BIBLIOGRAPHY

- [1] R. K. Brayton and J. K. Moser, A Theory of Nonlinear Networks.
IBM Research Paper RC-906, March 20, 1963, Yorktown Heights, N. Y.
- [2] Alfred North Whitehead, A Treatise on Universal Algebra, Hafner
Publishing Company, New York, 1960.
- [3] J. LaSalle and S. Lifschitz, Stability by Liapunov's Direct
Methods, Academic Press, New York, 1961, p. 66.