

CHAPTER 1

INTRODUCTORY CONCEPTS AND MATHEMATICS

PART I INTRODUCTION

1-1 Trends and Scopes

In the 21st century, the transcendent and translational technologies include nanotechnology, microelectronics, information technology, and biotechnology as well as the enabling and supporting mechanical and civil infrastructure systems and smart materials. These technologies are the primary drivers of the century and the new economy in a modern society. Mechanics forms the backbone and basis of these transcendent and translational technologies (Chong, 2004, 2010). Papers on the applications of the theory of elasticity to engineering problems form a significant part of the technical literature in solid mechanics (e.g. Dvorak, 1999; Oden, 2006). Many of the solutions presented in current papers employ numerical methods and require the use of high-speed digital computers. This trend is expected to continue into the foreseeable future, particularly with the widespread use of microcomputers and minicomputers as well as the increased availability of supercomputers (Londer, 1985; Fosdick, 1996). For example, finite element methods have been applied to a wide range of problems such as plane problems, problems of plates and shells, and general three-dimensional problems, including linear and nonlinear behavior, and isotropic and anisotropic materials. Furthermore, through the use of computers, engineers have been able to consider the optimization of large engineering systems (Atrek et al., 1984; Zienkiewicz and Taylor, 2005; Kirsch, 1993; Tsompanakis et al., 2008) such as the space shuttle. In addition, computers have played a powerful role

in the fields of computer-aided design (CAD) and computer-aided manufacturing (CAM) (Ellis and Semenov, 1983; Lamit, 2007) as well as in virtual testing and simulation-based engineering science (Fosdick, 1996; Yang and Pan, 2004; Oden, 2000, 2006).

At the request of one of the authors (Chong), Moon et al. (2003) conducted an in-depth National Science Foundation (NSF) workshop on the research needs of solid mechanics. The following are the recommendations.

Unranked overall priorities in solid mechanics research (Moon et al., 2003)

1. Modeling multiscale problems:
 - (i) Bridging the micro-nano-molecular scale
 - (ii) Macroscale dynamics of complex machines and systems
2. New experimental methods:
 - (i) Micro-nano-atomic scales
 - (ii) Coupling between new physical phenomena and model simulations
3. Micro- and nanomechanics:
 - (i) Constitutive models of failure initiation and evolution
 - (ii) Biocell mechanics
 - (iii) Force measurements in the nano- to femtonewton regime
4. Tribology, contact mechanics:
 - (i) Search for a grand theory of friction and adhesion
 - (ii) Molecular-atomic-based models
 - (iii) Extension of microscale models to macroapplications
5. Smart, active, self-diagnosis and self-healing materials:
 - (i) Microelectromechanical systems (MEMS)/Nanoelectromechanical systems (NEMS) and biomaterials
 - (ii) Fundamental models
 - (iii) Increased actuator capability
 - (iv) Application to large-scale devices and systems
6. Nucleation of cracks and other defects:
 - (i) Electronic materials
 - (ii) Nanomaterials
7. Optimization methods in solid mechanics:
 - (i) Synthesis of materials by design
 - (ii) Electronic materials
 - (iii) Optimum design of biomaterials
8. Nonclassical materials:
 - (i) Foams, granular materials, nanocarbon tubes, smart materials
9. Energy-related solid mechanics:
 - (i) High-temperature materials and coatings
 - (ii) Fuel cells

10. Advanced material processing:
 - (i) High-speed machining
 - (ii) Electronic and nanodevices, biodevices, biomaterials
11. Education in mechanics:
 - (i) Need for multidisciplinary education between solid mechanics, physics, chemistry, and biology
 - (ii) New mathematical skills in statistical mechanics and optimization methodology
12. Problems related to Homeland Security (Postworkshop; added by the editor)
 - (i) Ability of infrastructure to withstand destructive attacks
 - (ii) New safety technology for civilian aircraft
 - (iii) New sensors and robotics
 - (iv) New coatings for fire-resistant structures
 - (v) New biochemical filters

In addition to finite element methods, older techniques such as finite difference methods have also found applications in elasticity problems. More generally, the broad subject of approximation methods has received considerable attention in the field of elasticity. In particular, the boundary element method has been widely applied because of certain advantages it possesses in two- and three-dimensional problems and in infinite domain problems (Brebbia, 1988). In addition, other variations of the finite element method have been employed because of their efficiency. For example, finite strip, finite layer, and finite prism methods (Cheung and Tham, 1997) have been used for rectangular regions, and finite strip methods have been applied to nonrectangular regions by Yang and Chong (1984). This increased interest in approximate methods is due mainly to the enhanced capabilities of both mainframe and personal digital computers and their widespread use. Because this development will undoubtedly continue, the authors (Boresi, Chong, and Saigal) treat the topic of approximation methods in elasticity in a second book (Boresi et al., 2002), with particular emphasis on numerical stress analysis through the use of finite differences and finite elements, as well as boundary element and meshless methods.

However, in spite of the widespread use of approximate methods in elasticity (Boresi et al., 2002), the basic concepts of elasticity are fundamental and remain essential for the understanding and interpretation of numerical stress analysis. Accordingly, the present book devotes attention to the theories of deformation and of stress, the stress–strain relations (constitutive relations), nano- and biomechanics, and the fundamental boundary value problems of elasticity. Extensive use of index notation is made. However, general tensor notation is used sparingly, primarily in appendices.

In recent years, researchers from mechanics and other diverse disciplines have been drawn into vigorous efforts to develop smart or intelligent structures that can monitor their own condition, detect impending failure, control damage, and adapt

to changing environments (Rogers and Rogers, 1992). The potential applications of such smart materials/systems are abundant: design of smart aircraft skin embedded with fiber-optic sensors (Udd, 1995) to detect structural flaws, bridges with sensing/actuating elements to counter violent vibrations, flying microelectromechanical systems (Trimmer, 1990) with remote control for surveying and rescue missions, and stealth submarine vehicles with swimming muscles made of special polymers. Such a multidisciplinary infrastructural systems research front, represented by material scientists, physicists, chemists, biologists, and engineers of diverse fields—mechanical, electrical, civil, control, computer, aeronautical, and so on—has collectively created a new entity defined by the interface of these research elements. Smart structures/materials are generally created through synthesis by combining sensing, processing, and actuating elements integrated with conventional structural materials such as steel, concrete, or composites. Some of these structures/materials currently being researched or in use are listed below (Chong et al., 1990, 1994; Chong and Davis, 2000):

- Piezoelectric composites, which convert electric current to (or from) mechanical forces
- Shape memory alloys, which can generate force through changing the temperature across a transition state
- Electrorheological (ER) and magnetorheological (MR) fluids, which can change from liquid to solid (or the reverse) in electric and magnetic fields, respectively, altering basic material properties dramatically
- Bio-inspired sensors and nanotechnologies, e.g., graphenes and nanotubes

The science and technology of nanometer-scale materials, nanostructure-based devices, and their applications in numerous areas, such as functionally graded materials, molecular-electronics, quantum computers, sensors, molecular machines, and drug delivery systems—to name just a few, form the realm of *nanotechnology* (Srivastava et al., 2007). At nanometer length scale, the material systems concerned may be downsized to reach the limit of tens to hundreds of atoms, where many new physical phenomena are being discovered. Modeling of nanomaterials involving phenomena with multiple length/time scales has attracted enormous attention from the scientific research community. This is evidenced in the works of Belytschko et al. (2002), Belytschko and Xiao (2003), Liu et al. (2004), Arroyo and Belytschko (2005), Srivastava et al. (2007), Wagner et al. (2008), Masud and Kannan (2009), and the host of references mentioned therein. As a matter of fact, the traditional material models based on continuum descriptions are inadequate at the nanoscale, even at the microscale. Therefore, simulation techniques based on descriptions at the atomic scale, such as molecular dynamics (MD), has become an increasingly important computational toolbox. However, MD simulations on even the largest supercomputers (Abraham et al., 2002), although enough for the study of some nanoscale phenomena, are still far too small to treat the micro-to-macroscale interactions that must be captured in the simulation of any real device (Wagner et al., 2008).

Bioscience and technology has contributed much to our understanding of human health since the birth of continuum biomechanics in the mid-1960s (Fung, 1967, 1983, 1990, 1993, 1995). Nevertheless, it has yet to reach its full potential as a consistent contributor to the improvement of health-care delivery. This is due to the fact that most biological materials are very complicated hierarchical structures. In the most recent review paper, Meyers et al. (2008) describe the defining characteristics, namely, hierarchy, multifunctionality, self-healing, and self-organization of biological tissues in detail, and point out that the new frontiers of material and structure design reside in the synthesis of bioinspired materials, which involve nanoscale self-assembly of the components and the development of hierarchical structures. For example the amazing multiscale bones structure—from amino acids, tropocollagen, mineralized collagen fibrils, fibril arrays, fiber patterns, osteon and Haversian canal, and bone tissue to macroscopic bone—makes bones remarkably resistant to fracture (Ritchie et al., 2009). The multiscale bone structure of trabecular bone and cortical bone from nanoscale to macroscale is illustrated in Figure 1-1.1. (Courtesy of I. Jasiuk and E. Hamed, University of Illinois – Urbana). Although much significant progress has been made in the field of bioscience and technology, especially in biomechanics, there exist many open problems related to elasticity, including molecular and cell biomechanics, biomechanics of development, biomechanics of growth and remodeling, injury biomechanics and rehabilitation, functional tissue engineering, muscle mechanics and active stress, solid–fluid interactions, and thermal treatment (Humphrey, 2002).

Current research activities aim at understanding, synthesizing, and processing material systems that behave like biological systems. Smart structures/materials basically possess their own sensors (nervous system), processor (brain system), and actuators (muscular systems), thus mimicking biological systems (Rogers and Rogers, 1992). Sensors used in smart structures/materials include optical fibers, micro-cantilevers, corrosion sensors, and other environmental sensors and sensing particles. Examples of actuators include shape memory alloys that would return to their original shape when heated, hydraulic systems, and piezoelectric ceramic polymer composites. The processor or control aspects of smart structures/materials are based on microchip, computer software, and hardware systems.

Recently, Huang from Northwestern University and his collaborators developed the stretchable silicon based on the wrinkling of the thin films on a prestretched substrate. This is important to the development of stretchable electronics and sensors such as the three-dimensional eye-shaped sensors. One of their papers was published in *Science* in 2006 (Khang et al., 2006). The basic idea is to make straight silicon ribbons wavy. A prestretched polymer Polydimethylsiloxane (PDMS) is used to peel silicon ribbons away from the substrate, and releasing the prestretch leads to buckled, wavy silicon ribbons.

In the past, engineers and material scientists have been involved extensively with the characterization of given materials. With the availability of advanced computing, along with new developments in material sciences, researchers can now characterize processes, design, and manufacture materials with desirable performance and properties. Using nanotechnology (Reed and Kirk, 1989; Timp, 1999;

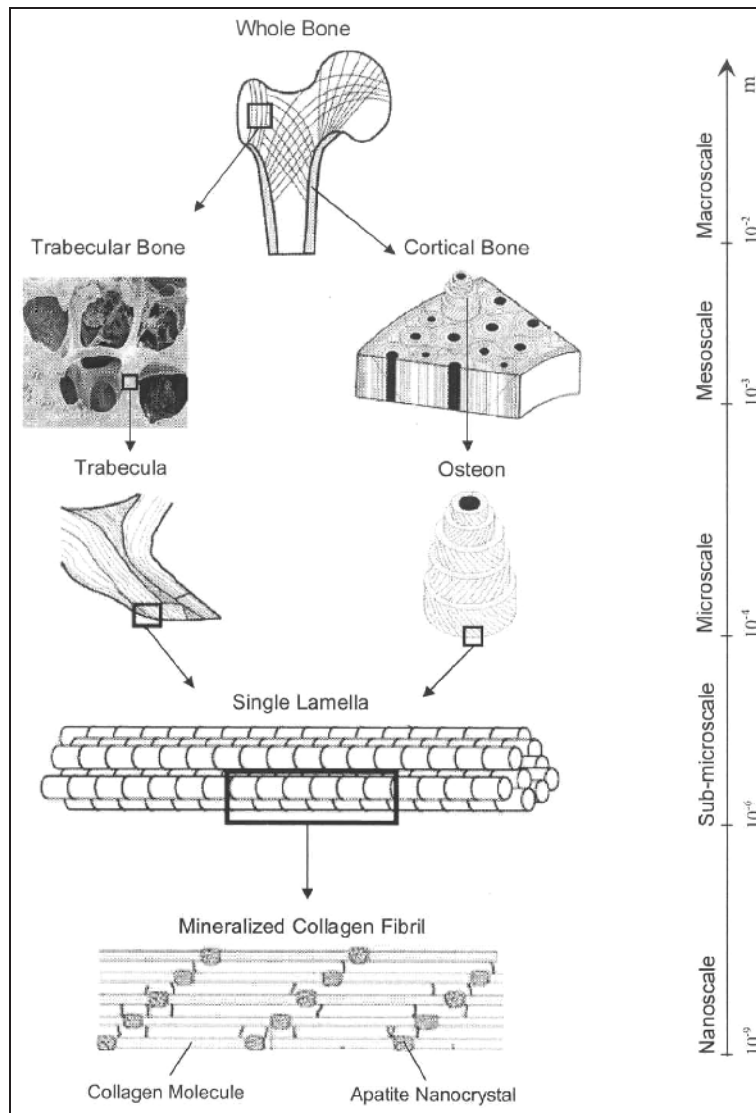


Figure 1-1.1

Chong, 2004), engineers and scientists can build designer materials molecule by molecule via self-assembly, etc. One of the challenges is to model short-term microscale material behavior through mesoscale and macroscale behavior into long-term structural systems performance (Fig. 1-1.2). Accelerated tests to simulate various environmental forces and impacts are needed. Supercomputers and/or workstations used in parallel are useful tools to (a) solve this multiscale and size-effect problem by taking into account the large number of variables and unknowns

MATERIALS		STRUCTURES		INFRASTRUCTURE
Nanolevel ~	microlevel ~	mesolevel ~	macro-level ~	systems integration
<i>Molecular Scale</i>	<i>Microns</i>		<i>Meters</i>	<i>Up to km Scale</i>
nanomechanics	micromechanics	mesomechanics	beams	bridge systems
self-assembly	microstructures	interfacial structures	columns	lifelines
nanofabrication	smart materials	composites	plates	airplanes

Figure 1-1.2 Scales in materials and structures.

to project microbehavior into infrastructure systems performance and (b) to model or extrapolate short-term test results into long-term life-cycle behavior.

According to Eugene Wong, the former engineering director of the National Science Foundation, the transcendent technologies of our time are

- Microelectronics—Moore's law: doubling the capabilities every 2 years for the past 30 years; unlimited scalability
- Information technology: confluence of computing and communications
- Biotechnology: molecular secrets of life

These technologies and nanotechnology are mainly responsible for the tremendous economic developments. Engineering mechanics is related to all these technologies based on the experience of the authors. The first small step in many of these research activities and technologies involves the study of deformation and stress in materials, along with the associated stress–strain relations.

In this book following the example of modern continuum mechanics and the example of A. E. Love (Love, 2009), we treat the theories of deformation and of stress separately, in this manner clearly noting their mathematical similarities and their physical differences. Continuum mechanics concepts such as couple stress and body couple are introduced into the theory of stress in the appendices of Chapters 3, 5, and 6. These effects are introduced into the theory in a direct way and present no particular problem. The notations of stress and of strain are based on the concept of a continuum, that is, a continuous distribution of matter in the region (space) of interest. In the mathematical physics sense, this means that the volume or region under examination is sufficiently filled with matter (dense) that concepts such as mass density, momentum, stress, energy, and so forth are defined at all points in the region by appropriate mathematical limiting processes (see Chapter 3, Section 3-1).

1-2 Theory of Elasticity

The theory of elasticity, in contrast to the general theory of continuum mechanics (Eringen, 1980), is an ad hoc theory designed to treat explicitly a special response

of materials to applied forces—namely, the elastic response, in which the stress at every point P in a material body (continuum) depends at all times solely on the simultaneous deformation in the immediate neighborhood of the point P (see Chapter 4, Section 4-1). In general, the relation between stress and deformation is a nonlinear one, and the corresponding theory is called the *nonlinear theory of elasticity* (Green and Adkins, 1970). However, if the relationship of the stress and the deformation is linear, the material is said to be linearly elastic, and the corresponding theory is called the *linear theory of elasticity*.

The major part of this book treats the linear theory of elasticity. Although ad hoc in form, this theory of elasticity plays an important conceptual role in the study of nonelastic types of material responses. For example, often in problems involving plasticity or creep of materials, the method of successive elastic solutions is employed (Mendelson, 1983). Consequently, the theory of elasticity finds application in fields that treat inelastic response.

1-3 Numerical Stress Analysis

The solution of an elasticity problem generally requires the description of the response of a material body (computer chips, machine part, structural element, or mechanical system) to a given excitation (such as force). In an engineering sense, this description is usually required in numerical form, the objective being to assure the designer or engineer that the response of the system will not violate design requirements. These requirements may include the consideration of deterministic and probabilistic concepts (Thoft-Christensen and Baker, 1982; Wen, 1984; Yao, 1985). In a broad sense the numerical results are predictions as to whether the system will perform as desired. The solution to the elasticity problem may be obtained by a direct numerical process (numerical stress analysis) or in the form of a general solution (which ordinarily requires further numerical evaluation; see Section 1-4).

The usual methods of numerical stress analysis recast the mathematically posed elasticity problem into a direct numerical analysis. For example, in finite difference methods, derivatives are approximated by algebraic expressions; this transforms the differential boundary value problem of elasticity into an algebraic boundary value problem requiring the numerical solution of a set of simultaneous algebraic equations. In finite element methods, trial function approximations of displacement components, stress components, and so on are employed in conjunction with energy methods (Chapter 4, Section 4-21) and matrix methods (Section 1-28), again to transform the elasticity boundary value problem into a system of simultaneous algebraic equations. However, because finite element methods may be applied to individual pieces (elements) of the body, each element may be given distinct material properties, thus achieving very general descriptions of a body as a whole. This feature of the finite element method is very attractive to the practicing stress analyst. In addition, the application of finite elements leads to many interesting mathematical questions concerning accuracy of approximation, convergence of the results, attainment of bounds on the exact answer, and so on. Today, finite element

methods are perhaps the principal method of numerical stress analysis employed to solve elasticity problems in engineering (Zienkiewicz and Taylor, 2005). By their nature, methods of numerical stress analysis (Boresi et al., 2002) yield approximate solutions to the exact elasticity solution.

1-4 General Solution of the Elasticity Problem

Plane Elasticity. Two classical plane problems have been studied extensively: plane strain and plane stress (see Chapter 5). If the state of plane isotropic elasticity is referred to the (x, y) plane, then plane elasticity is characterized by the conditions that the stress and strain are independent of coordinate z , and shear stress τ_{xz} , τ_{yz} (hence, shear strains γ_{xz} , γ_{yz}) are zero. In addition, for plane strain the extensional strain ϵ_z equals 0, and for plane stress we have $\sigma_z = 0$. For plane strain problems the equations represent exact solutions to physical problems, whereas for plane stress problems, the usual solutions are only approximations to physical problems. Mathematically, the problems of plane stress and plane strain are identical (see Chapter 5).

One general method of solution of the plane problem rests on the reduction of the elasticity equations to the solution of certain equations in the complex plane (Muskhelishvili, 1975).¹ Ordinarily, the method requires mapping of the given region into a suitable region in the complex plane. A second general method rests on the introduction of a single scalar biharmonic function, the Airy stress function, which must be chosen suitably to satisfy boundary conditions (see Chapter 5).

Three-Dimensional Elasticity. In contrast to the problem of plane elasticity, the construction of general solutions of the three-dimensional equations of elasticity has not as yet been completely achieved. Many so-called general solutions are really particular forms of solutions of the three-dimensional field equations of elasticity in terms of arbitrary, ad hoc functions. Particular examples of general solutions are employed in Chapter 8 and in Appendix 5B. In many of these examples, the functions and the form of solution are determined in part by the differential equations and in part by the physical features of the problem. A general solution of the elasticity equations may also be constructed in terms of biharmonic functions (see Appendix 5B). Because there is no apparent reason for one form of general solution to be readily obtainable from another, a number of investigators have attempted to extend the generality of solution form and show relations among known solutions (Sternberg, 1960; Naghdi and Hsu, 1961; Stippes, 1967).

1-5 Experimental Stress Analysis

Material properties that enter into the stress–strain relations (constitutive relations; see Section 4-4) must be obtained experimentally (Schreiber et al., 1973; Chong and Smith, 1984). In addition, other material properties, such as ultimate strength

¹See also Appendix 5B.

and fracture toughness, as well as nonmaterial quantities such as residual stresses, have to be determined by physical tests.

For bodies that possess intricately shaped boundaries, general analytical (closed-form) solutions become extremely difficult to obtain. In such cases one must invariably resort to approximate methods, principally to numerical methods or to experimental methods. In the latter, several techniques such as photoelasticity, the Moiré method, strain gage methods, fracture gages, optical fibers, and so forth have been developed to a fine art (Dove and Adams, 1964; Dally and Riley, 2005; Rogers and Rogers, 1992; Ruud and Green, 1984). In addition, certain analogies based on a similarity between the equations of elasticity and the equations that describe readily studied physical systems are employed to obtain estimates of solutions or to gain insight into the nature of mathematical solutions (see Chapter 7, Section 7-9, for the membrane analogy in torsion). In this book we do not treat experimental methods but rather refer to the extensive modern literature available.²

1-6 Boundary Value Problems of Elasticity

The solution of the equations of elasticity involves the determination of a stress or strain state in the interior of a region R subject to a given state of stress or strain (or displacement) on the boundary B of R (see Chapter 4, Section 4-15). Subject to certain restrictions on the nature of the solution and of region R and the form of the boundary conditions, the solution of boundary value problems of elasticity may be shown to exist (see Chapter 4, Section 4-16). Under broader conditions, existence and uniqueness of the elasticity boundary value problem are not ensured. In general, the question of existence and uniqueness (Knops and Payne, 1971) rests on the theory of systems of partial differential equations of three independent variables.

In particular forms the boundary value problem of elasticity may be reduced to that of seeking a single scalar function f of three independent variables, say (x, y, z) ; that is, $f = f(x, y, z)$ such that the stress field or strain field derived from f satisfies the boundary conditions on B . In particular for the Laplace equation, three types of boundary value problems occur frequently in elasticity: the Dirichlet problem, the Neumann problem, and the mixed problem. Let $h(x, y)$ be a given function that is defined on B , the bounding surface of a simply connected region R . Then the Dirichlet problem for the Laplace equation is that of determining a function $f = f(x, y)$ that

1. is continuous on $R + B$,
2. is harmonic on R , and
3. is identical to $h(x, y)$ on B .

²*Experimental Mechanics* and *Experimental Techniques*, both journals of the Society for Experimental Mechanics (SEM), contain a wealth of information on experimental techniques. In addition, the American Society for Testing and Materials (ASTM) publishes the *Journal of Testing and Evaluation*, the *Geotechnical Testing Journal*, and other journals.

The Dirichlet problem has been shown to possess a unique solution (Greenspan, 1965). However, analytical determination of $f(x, y)$ is very much more difficult to achieve than is the establishment of its existence. Indeed, except for special forms of boundary B (such as the rectangle, the circle, or regions that can be mapped onto rectangular or circular regions), the problems of determining $f(x, y)$ do not surrender to existing analytical techniques.

The Neumann boundary value problem for the Laplace equation is that of determining a function $f(x, y)$ that

1. is defined and continuous on $R + B$,
2. is harmonic on R , and
3. has an outwardly directed normal derivative $\partial f/\partial n$ such that $\partial f/\partial n = g(x, y)$ on B , where $g(x, y)$ is defined and continuous on B .

Without an additional requirement [namely, that $f(x, y)$ has a prescribed value for at least one point of B], the solution of the Neumann problem is not well posed because otherwise the Neumann problem has a one-parameter infinity of solutions.

The mixed problem overcomes the difficulty of the Neumann problem. Again, let $g(x, y)$ be a continuous function on B' of R and let $h(x, y)$ be bounded and continuous on B'' of R , where $B = B' + B''$ denotes the boundary of region R . Then the mixed problem for the Laplace equation is that of determining a function $f(x, y)$ such that it

1. is defined and continuous on $R + B$,
2. is harmonic on R ,
3. is identical with $g(x, y)$, on B' , and
4. has outwardly directed normal derivative $\partial f/\partial n = h(x, y)$ on B'' .

It has been shown that certain mixed problems have unique solutions³ (Greenspan, 1965). Because, in general, the solutions of the Dirichlet and mixed problems cannot be given in closed form, methods of approximate solutions of these problems are presented in another book by the authors (Boresi et al., 2002). More generally, these approximate methods may be applied to most boundary value problems of elasticity.

PART II PRELIMINARY CONCEPTS

In Part II of this chapter we set down some concepts that are useful in following the developments in the text proper and in the appendices.

³These remarks are restricted to simply connected regions.

1-7 Brief Summary of Vector Algebra

In this text a boldface letter denotes a vector quantity unless an explicit statement to the contrary is given; thus, $\mathbf{A}$ denotes a vector. Frequently, we denote a vector by the set of its projections (A_x, A_y, A_z) on rectangular Cartesian axes (x, y, z) . Thus,

$$\mathbf{A} = (A_x, A_y, A_z) \quad (1-7.1)$$

The magnitude of a vector $\mathbf{A}$ is denoted by

$$|\mathbf{A}| = A = (A_x^2 + A_y^2 + A_z^2)^{1/2} \quad (1-7.2)$$

We may also express a vector in terms of its components with respect to (x, y, z) axes. For example,

$$\mathbf{A} = \mathbf{i}A_x + \mathbf{j}A_y + \mathbf{k}A_z \quad (1-7.3)$$

where $\mathbf{i}A_x, \mathbf{j}A_y, \mathbf{k}A_z$ are components of $\mathbf{A}$ with respect to axes (x, y, z) , and $\mathbf{i}, \mathbf{j}, \mathbf{k}$, are unit vectors directed along positive (x, y, z) axes, respectively. In general, the symbols $\mathbf{i}, \mathbf{j}, \mathbf{k}$ denote unit vectors.

Vector quantities obey the *associative law of vector addition*:

$$\mathbf{A} + (\mathbf{B} + \mathbf{C}) = (\mathbf{A} + \mathbf{B}) + \mathbf{C} = \mathbf{A} + \mathbf{B} + \mathbf{C} \quad (1-7.4)$$

and the *commutative law of vector addition*:

$$\mathbf{A} + \mathbf{B} = \mathbf{B} + \mathbf{A} \quad \mathbf{A} + \mathbf{B} + \mathbf{C} = \mathbf{B} + \mathbf{A} + \mathbf{C} = \mathbf{B} + \mathbf{C} + \mathbf{A} \quad (1-7.5)$$

Symbolically, we may represent a vector quantity by an arrow (Fig. 1-7.1) with the understanding that the addition of any two arrows (vectors) must obey the commutative law [Eq. (1-7.5)].

The *scalar product* of two vectors $\mathbf{A}, \mathbf{B}$ is defined to be

$$\mathbf{A} \cdot \mathbf{B} = A_x B_x + A_y B_y + A_z B_z \quad (1-7.6)$$

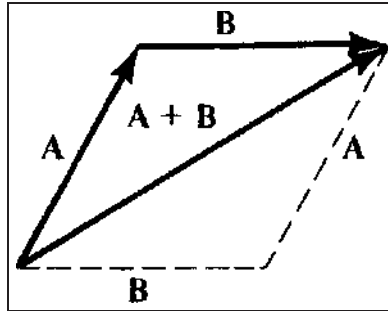


Figure 1-7.1

where the symbol $\cdot$ is a conventional notation for the scalar product. By the above definition, it follows that the scalar product of vectors is commutative; that is,

$$\mathbf{A} \cdot \mathbf{B} = \mathbf{B} \cdot \mathbf{A} \quad (1-7.7)$$

A useful property of the scalar product of two vectors is

$$\mathbf{A} \cdot \mathbf{B} = AB \cos \theta \quad (1-7.8)$$

where A and B denote the magnitudes of vectors $\mathbf{A}$ and $\mathbf{B}$, respectively, and the angle θ denotes the angle formed by vectors $\mathbf{A}$ and $\mathbf{B}$ (Fig. 1-7.2).

If $\mathbf{B}$ is a unit vector in the x direction, Eqs. (1-7.3) and (1-7.8) yield $A_x = A \cos \alpha$, where α is the direction angle between the vector $\mathbf{A}$ and the positive x axis. Similarly, $A_y = A \cos \beta$, $A_z = A \cos \gamma$, where β, γ denote direction angles between the vector $\mathbf{A}$ and the y axis and the z axis, respectively. Substitution of these expressions into Eq. (1-7.2) yields the relation

$$\cos^2 \alpha + \cos^2 \beta + \cos^2 \gamma = 1 \quad (1-7.9)$$

Thus, the direction cosines of vector $\mathbf{A}$ are *not independent*. They must satisfy Eq. (1-7.9).

The scalar product law of vectors has other properties in common with the product of numbers. For example,

$$\mathbf{A} \cdot (\mathbf{B} + \mathbf{C}) = \mathbf{A} \cdot \mathbf{B} + \mathbf{A} \cdot \mathbf{C} \quad (1-7.10)$$

$$\begin{aligned} (\mathbf{A} + \mathbf{B}) \cdot (\mathbf{C} + \mathbf{D}) &= (\mathbf{A} + \mathbf{B}) \cdot \mathbf{C} + (\mathbf{A} + \mathbf{B}) \cdot \mathbf{D} \\ &= \mathbf{A} \cdot \mathbf{C} + \mathbf{B} \cdot \mathbf{C} + \mathbf{A} \cdot \mathbf{D} + \mathbf{B} \cdot \mathbf{D} \end{aligned} \quad (1-7.11)$$

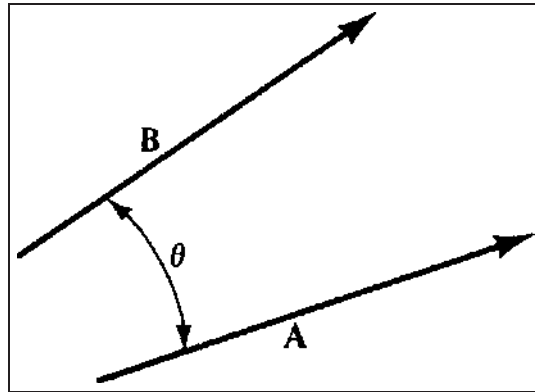


Figure 1-7.2

The *vector product* of two vectors **A** and **B** is defined to be a third vector **C** whose magnitude is given by the relation

$$C = AB \sin \theta \quad (1-7.12)$$

The direction of vector **C** is perpendicular to the plane formed by vectors **A** and **B**. The sense of **C** is such that the three vectors **A**, **B**, **C** form a right-handed or left-handed system according to whether the coordinate system (x, y, z) is right handed or left handed (see Fig. 1-7.3).

Symbolically, we denote the vector product of **A** and **B** in the form

$$\mathbf{C} = \mathbf{A} \times \mathbf{B} \quad (1-7.13)$$

where $\times$ denotes vector product (or cross product). In determinant notation, Eq. (1-7.13) may be written as

$$\mathbf{C} = \begin{vmatrix} \mathbf{i} & \mathbf{j} & \mathbf{k} \\ A_x & A_y & A_z \\ B_x & B_y & B_z \end{vmatrix} \quad (1-7.13a)$$

where $(A_x, A_y, A_z), (B_x, B_y, B_z)$ denotes $(\mathbf{i}, \mathbf{j}, \mathbf{k})$ projections of vectors **(A, B)**, respectively.

The vector product of vectors has the following property:

$$\mathbf{A} \times \mathbf{B} = -\mathbf{B} \times \mathbf{A} \quad (1-7.14)$$

Accordingly, the vector product of vectors is not commutative.

The vector product also has the following properties:

$$\begin{aligned} \mathbf{R} \times (\mathbf{A} + \mathbf{B}) &= \mathbf{R} \times \mathbf{A} + \mathbf{R} \times \mathbf{B} \\ (\mathbf{A} + \mathbf{B}) \times \mathbf{R} &= \mathbf{A} \times \mathbf{R} + \mathbf{B} \times \mathbf{R} \end{aligned} \quad (1-7.15)$$

$$\begin{aligned} (\mathbf{A} + \mathbf{B}) \times (\mathbf{C} + \mathbf{D}) &= (\mathbf{A} + \mathbf{B}) \times \mathbf{C} + (\mathbf{A} + \mathbf{B}) \times \mathbf{D} \\ &= \mathbf{A} \times \mathbf{C} + \mathbf{B} \times \mathbf{C} + \mathbf{A} \times \mathbf{D} + \mathbf{B} \times \mathbf{D} \end{aligned} \quad (1-7.16)$$

The scalar triple product of three vectors **A**, **B**, **C** is defined by the relation

$$\begin{aligned} \mathbf{A} \cdot (\mathbf{B} \times \mathbf{C}) &= A_x(B_y C_z - B_z C_y) + A_y(B_z C_x - B_x C_z) \\ &\quad + A_z(B_x C_y - B_y C_x) \end{aligned} \quad (1-7.17)$$

In determinant notation, the scalar triple product is

$$\mathbf{A} \cdot (\mathbf{B} \times \mathbf{C}) = \begin{vmatrix} A_x & A_y & A_z \\ B_x & B_y & B_z \\ C_x & C_y & C_z \end{vmatrix} \quad (1-7.18)$$

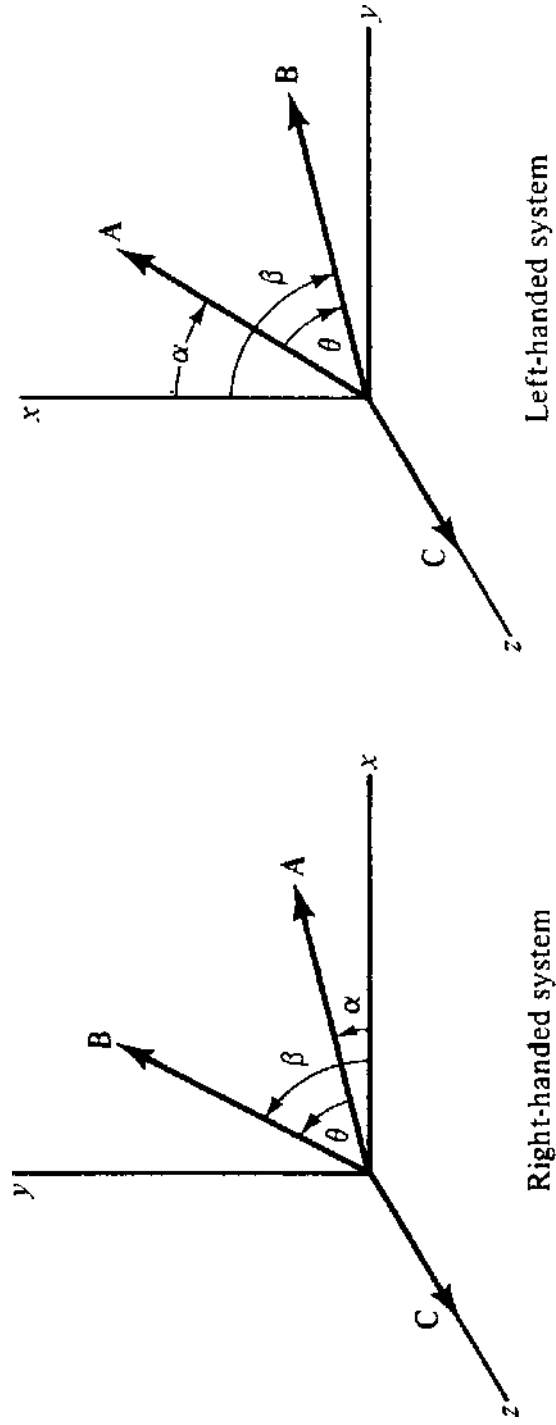


Figure 1-7.3

Because only the sign of a determinant changes when two rows are interchanged, two consecutive transpositions of rows leave a determinant unchanged. Consequently,

$$\mathbf{A} \cdot (\mathbf{B} \times \mathbf{C}) = \mathbf{C} \cdot (\mathbf{A} \times \mathbf{B}) = \mathbf{B} \cdot (\mathbf{C} \times \mathbf{A}) \quad (1-7.19)$$

Another useful property is the relation

$$(\mathbf{A} \times \mathbf{B}) \cdot \mathbf{C} = \mathbf{A} \cdot (\mathbf{B} \times \mathbf{C}) \quad (1-7.20)$$

The vector triple product of three vectors $\mathbf{A}$, $\mathbf{B}$, $\mathbf{C}$ is defined as

$$\mathbf{A} \times (\mathbf{B} \times \mathbf{C}) = \mathbf{B}(\mathbf{A} \cdot \mathbf{C}) - \mathbf{C}(\mathbf{A} \cdot \mathbf{B}) \quad (1-7.21)$$

Furthermore,

$$\begin{aligned} (\mathbf{A} \times \mathbf{B}) \cdot (\mathbf{C} \times \mathbf{D}) &= \mathbf{A} \cdot \mathbf{B} \times (\mathbf{C} \times \mathbf{D}) = \mathbf{A} \cdot [\mathbf{C}(\mathbf{B} \cdot \mathbf{D}) - (\mathbf{B} \cdot \mathbf{C})\mathbf{D}] \\ &= (\mathbf{A} \cdot \mathbf{C})(\mathbf{B} \cdot \mathbf{D}) - (\mathbf{A} \cdot \mathbf{D})(\mathbf{B} \cdot \mathbf{C}) \end{aligned} \quad (1-7.22)$$

Equation (1-7.22) follows from Eqs. (1-7.20) and (1-7.21).

1-8 Scalar Point Functions

Any scalar function $f(x, y, z)$ that is defined at all points in a region of space is called a *scalar point function*. Conceivably, the function f may depend on time, but if it does, attention can be confined to conditions at a particular instant. The region of space in which f is defined is called a scalar field. It is assumed that f is differentiable in this *scalar field*. Physical examples of scalar point functions are the mass density of a compressible medium, the temperature in a body, the flux density in a nuclear reactor, and the potential in an electrostatic field.

Consider the rate of change of the function f in various directions at some point P : (x, y, z) in the scalar field for which f is defined. Let (x, y, z) take increments (dx, dy, dz) . Then the function f takes an increment:

$$df = \frac{\partial f}{\partial x}dx + \frac{\partial f}{\partial y}dy + \frac{\partial f}{\partial z}dz \quad (1-8.1)$$

Consider the infinitesimal vector $\mathbf{i}dx + \mathbf{j}dy + \mathbf{k}dz$, where $(\mathbf{i}, \mathbf{j}, \mathbf{k})$ are unit vectors in the (x, y, z) directions, respectively. Its magnitude is $ds = (dx^2 + dy^2 + dz^2)^{1/2}$, and its direction cosines are

$$\cos \alpha = \frac{dx}{ds} \quad \cos \beta = \frac{dy}{ds} \quad \cos \gamma = \frac{dz}{ds}$$

The vector $\mathbf{i}(dx/ds) + \mathbf{j}(dy/ds) + \mathbf{k}(dz/ds)$ is a unit vector in the direction of $\mathbf{i}dx + \mathbf{j}dy + \mathbf{k}dz$, as division of a vector by a scalar alters only the magnitude of

the vector. Dividing Eq. (1-8.1) by ds , we obtain

$$\frac{df}{ds} = \frac{\partial f}{\partial x} \frac{dx}{ds} + \frac{\partial f}{\partial y} \frac{dy}{ds} + \frac{\partial f}{\partial z} \frac{dz}{ds}$$

or

$$\frac{df}{ds} = \frac{\partial f}{\partial x} \cos \alpha + \frac{\partial f}{\partial y} \cos \beta + \frac{\partial f}{\partial z} \cos \gamma \quad (1-8.2)$$

From Eq. (1-8.2) it is apparent that df/ds depends on the direction of ds ; that is, it depends on the direction (α, β, γ) . For this reason df/ds is known as the *directional derivative* of f in the direction (α, β, γ) . It represents the rate of change of f in the direction (α, β, γ) . For example, if $\alpha = 0$, $\beta = \gamma = \pi/2$,

$$\frac{df}{ds} = \frac{\partial f}{\partial x}$$

This is the rate of change of f in the direction of the x axis.

Maximum Value of the Directional Derivative. Gradient. By definition of the scalar product of two vectors, Eq. (1-8.2) may be written in the form

$$\frac{df}{ds} = \mathbf{n} \cdot \text{grad } f \quad (1-8.3)$$

where $\mathbf{n} = \mathbf{i} \cos \alpha + \mathbf{j} \cos \beta + \mathbf{k} \cos \gamma$ is a unit vector in the direction (α, β, γ) , and

$$\text{grad } f = \mathbf{i} \frac{\partial f}{\partial x} + \mathbf{j} \frac{\partial f}{\partial y} + \mathbf{k} \frac{\partial f}{\partial z} \quad (1-8.4)$$

is a vector point function (see Section 1-10) of (x, y, z) called the *gradient* of the scalar function f . Because $\mathbf{n}$ is a unit vector, Eq. (1-8.3) shows that $|\text{grad } f|$ is the maximum value of df/ds at the point $P: (x, y, z)$ and that the direction of $\text{grad } f$ is the direction in which $f(x, y, z)$ *increases* most rapidly. Equation (1-8.3) also shows that the directional derivative of f in any direction is the component of the vector $\text{grad } f$ in that direction.

The equation $f(x, y, z) = C$ defines a family of surfaces, one surface for each value of the constant C . These are called *level surfaces* of the function f . If $\mathbf{n}$ is tangent to a level surface, the directional derivative of f in the direction of $\mathbf{n}$ is zero, as f is constant along a *level surface*. Consequently, by Eq. (1-8.3), the vector $\mathbf{n}$ must be perpendicular to the vector $\text{grad } f$ when $\mathbf{n}$ is tangent to a level surface. Accordingly, the vector $\text{grad } f$ at the point $P: (x, y, z)$ is normal to the level surface of f through the point $P: (x, y, z)$.

A symbolic vector operator, called *del* or *nabla*, is defined as follows:

$$\nabla = \mathbf{i} \frac{\partial}{\partial x} + \mathbf{j} \frac{\partial}{\partial y} + \mathbf{k} \frac{\partial}{\partial z} \quad (1-8.5)$$

By Eqs. (1-8.3), (1-8.4), and (1-8.5),

$$\text{grad } f = \nabla f$$

and

$$\frac{df}{ds} = \mathbf{n} \cdot \nabla f$$

By definition,

$$\nabla \cdot \nabla = \nabla^2 = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2} \quad (1-8.6)$$

Consequently, the Laplace equation may be written symbolically as

$$\nabla^2 f = \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2} + \frac{\partial^2 f}{\partial z^2} = 0 \quad (1-8.7)$$

For this reason the symbolic operator ∇^2 is called the *Laplacian*.

1-9 Vector Fields

Assume that for each point $P: (x, y, z)$ in a region there exists a vector point function $\mathbf{q}(x, y, z)$. This vector point function is called a *vector field*. It may be represented at each point in the region by a vector with length equal to the magnitude of $\mathbf{q}$ and drawn in the direction of $\mathbf{q}$. For example, for each point in a flowing fluid there corresponds a vector $\mathbf{q}$ that represents the velocity of the particle of fluid at that point. This vector point function is called the velocity field of the fluid. Another example of a vector field is the displacement vector function for the particles of a deformable body. Electric and magnetic field intensities are also vector fields. A vector field is often simply called a “vector.”

In any continuous vector field there exists a system of curves such that the vectors along a curve are everywhere tangent to the curve; that is, the vector field consists exclusively of tangent vectors to the curves. These curves are called the *vector lines* (or *field lines*) of the field. The vector lines of a velocity field are called *stream lines*. The vector lines in an electrostatic or magnetostatic field are known as *lines of force*. In general, the vector function $\mathbf{q}$ may depend on (x, y, z) and t , where t denotes time. If $\mathbf{q}$ depends on time, the field is said to be *unsteady* or *nonstationary*; that is, the field varies with time. For a *steady field*, $\mathbf{q} = \mathbf{q}(x, y, z)$. For example, if a velocity field changes with time (i.e., if the flow is unsteady), the stream lines may change with time.

A vector field $\mathbf{q} = \mathbf{i}u + \mathbf{j}v + \mathbf{k}w$ is defined by expressing the projections (u, v, w) as functions of (x, y, z) . If (dx, dy, dz) is an infinitesimal vector in the direction of the vector $\mathbf{q}$, the direction cosines of this vector are $dx/ds = u/q$, $dy/ds = v/q$, and $dz/ds = w/q$. Consequently, the differential equations of the system of vector lines of the field are

$$\frac{ds}{q} = \frac{dx}{u} = \frac{dy}{v} = \frac{dz}{w} \quad (1-9.1)$$

In Eq. (1-9.1) the components (u, v, w) are functions of (x, y, z) . The finite equations of the system of vector lines are obtained by integrating Eq. (1-9.1). The theory of integration of differential equations of this type is explained in most books on differential equations (Morris and Brown, 1964; Ince, 2009).

If a given vector field $\mathbf{q}$ is the gradient of a scalar field f (i.e., if $\mathbf{q} = \text{grad } f$), the scalar function f is called a potential function for the vector field, and the vector field is called a potential field. Because $\text{grad } f$ is perpendicular to the level surfaces of f , it follows that the vector lines of a potential field are everywhere normal to the level surfaces of the potential function.

1-10 Differentiation of Vectors

An infinitesimal increment $d\mathbf{R}$ of a vector $\mathbf{R}$ need not be collinear with the vector $\mathbf{R}$ (Fig. 1-10.1). Consequently, in general, the vector $\mathbf{R} + d\mathbf{R}$ differs from the vector $\mathbf{R}$ not only in magnitude but also in direction. It would be misleading to denote the magnitude of the vector $d\mathbf{R}$ by dR , as dR denotes the increment of the magnitude R . Accordingly, the magnitude of $d\mathbf{R}$ is denoted by $|d\mathbf{R}|$ or by another symbol, such as ds . The magnitude of the vector $\mathbf{R} + d\mathbf{R}$ is $R + dR$. Figure 1-10.1 shows that $|\mathbf{R} + d\mathbf{R}| \leq R + |d\mathbf{R}|$. Hence, $dR \leq |d\mathbf{R}|$.

If the vector $\mathbf{R}$ is a function of a scalar t (where t may or may not denote time), $d\mathbf{R}/dt$ is defined to be a vector in the direction of $d\mathbf{R}$, with magnitude ds/dt (where $ds = |d\mathbf{R}|$).

Vectors obey the same rules of differentiation as scalars. This fact may be demonstrated by the Δ method that is used for deriving differentiation formulas

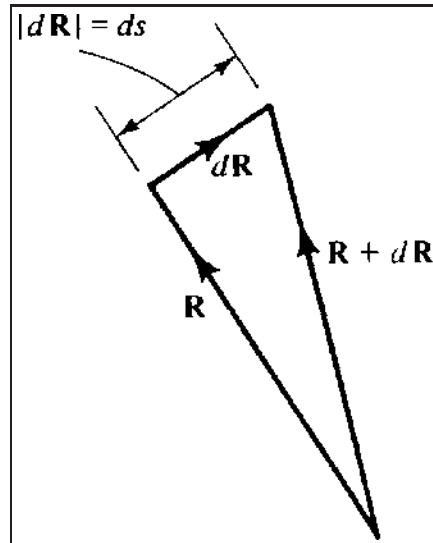


Figure 1-10.1

in scalar calculus. For example, consider the derivative of the vector function $\mathbf{Q} = u\mathbf{R}$, where u is a scalar function of t and $\mathbf{R}$ is a vector function of t . If t takes an increment Δt , $\mathbf{R}$ and u take increments $\Delta\mathbf{R}$ and Δu . Hence,

$$\mathbf{Q} + \Delta\mathbf{Q} = (u + \Delta u)(\mathbf{R} + \Delta\mathbf{R})$$

Subtracting $\mathbf{Q} = u\mathbf{R}$ and dividing by Δt , we obtain

$$\frac{\Delta\mathbf{Q}}{\Delta t} = \mathbf{R} \frac{\Delta u}{\Delta t} + u \frac{\Delta\mathbf{R}}{\Delta t} + \Delta u \frac{\Delta\mathbf{R}}{\Delta t}$$

As $\Delta t \rightarrow 0$, $\Delta u \rightarrow 0$, $\Delta\mathbf{Q}/\Delta t \rightarrow d\mathbf{Q}/dt$, $\Delta u/\Delta t \rightarrow du/dt$, and $\Delta\mathbf{R}/\Delta t \rightarrow d\mathbf{R}/dt$. Hence,

$$\frac{d\mathbf{Q}}{dt} = \mathbf{R} \frac{du}{dt} + u \frac{d\mathbf{R}}{dt} \quad (1-10.1)$$

Equation (1-10.1) has the same form as the formula for the derivative of the product of two scalars.

Let $\mathbf{R} = i\mathbf{u} + \mathbf{j}v + \mathbf{k}w$ be a single vector (not a vector field) where (i, j, k) are unit vectors and (u, v, w) are the (i, j, k) projections of $\mathbf{R}$, respectively. Let (u, v, w) take increments (du, dv, dw) . Then because (i, j, k) are constants, $\mathbf{R}$ takes the increment $d\mathbf{R} = i du + \mathbf{j} dv + \mathbf{k} dw$ where, in general, $d\mathbf{R}$ is not collinear with $\mathbf{R}$. If (u, v, w) are functions of the single variable t ,

$$\frac{d\mathbf{R}}{dt} = i \frac{du}{dt} + \mathbf{j} \frac{dv}{dt} + \mathbf{k} \frac{dw}{dt} \quad (1-10.2)$$

Hence, $d\mathbf{R}/dt$ is a vector in the direction of $d\mathbf{R}$, with magnitude $[(du/dt)^2 + (dv/dt)^2 + (dw/dt)^2]^{1/2}$.

If $\mathbf{R}$ is the position of a moving particle P measured from a fixed point O (Fig. 1-10.2), $d\mathbf{R}/dt$ is the velocity vector $\mathbf{q}$ of the particle. Likewise,

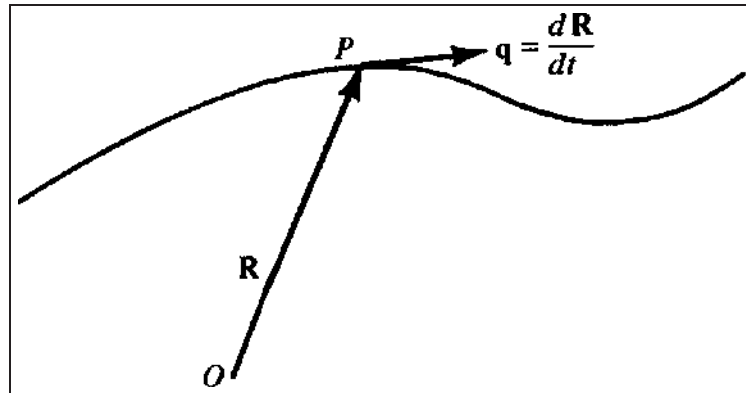


Figure 1-10.2

$d\mathbf{q}/dt = d^2\mathbf{R}/dt^2$ is the acceleration vector of the particle. Hence, the vector form of Newton's second law is

$$\mathbf{F} = m \frac{d^2\mathbf{R}}{dt^2} \quad (1-10.3)$$

1-11 Differentiation of a Scalar Field

Let $Q(x, y, z; t)$ be a scalar point function in a flowing fluid (such as temperature, density, a velocity projection, etc.). Then

$$dQ = \frac{\partial Q}{\partial x}dx + \frac{\partial Q}{\partial y}dy + \frac{\partial Q}{\partial z}dz + \frac{\partial Q}{\partial t}dt \quad (1-11.1)$$

Here (dx, dy, dz, dt) are arbitrary increments of coordinates (x, y, z) and time t . [In deformation theory, x, y, z are called *spatial (Eulerian) coordinates*; see Chapter 2.]

Let (dx, dy, dz) be the displacement that a particle of fluid experiences during a time interval dt . Then $dx/dt = u$, $dy/dt = v$, and $dz/dt = w$, where (u, v, w) is the velocity field. Hence, on dividing Eq. (1-11.1) by dt , we get

$$\frac{dQ}{dt} = u \frac{\partial Q}{\partial x} + v \frac{\partial Q}{\partial y} + w \frac{\partial Q}{\partial z} + \frac{\partial Q}{\partial t} \quad (1-11.2)$$

or, in vector notation,

$$\frac{dQ}{dt} = \mathbf{q} \cdot \text{grad } Q + \frac{\partial Q}{\partial t} \quad (1-11.3)$$

where $\mathbf{q}$ is the velocity field. Although Eq. (1-11.2) is derived for a scalar point function in a flowing fluid, it remains valid for any scalar point function $Q(x, y, z; t)$.

The distinction between $\partial Q/\partial t$ and dQ/dt is very important. The partial derivative $\partial Q/\partial t$ denotes the rate of change of Q at a fixed point of space as the fluid flows by. For steady flow, $\partial Q/\partial t = 0$. In contrast, dQ/dt denotes the rate of change of Q for a certain particle of fluid. For example, if Q is temperature, we determine $\partial Q/\partial t$ by holding the thermometer still. To determine dQ/dt , we must move the thermometer so that it coincides continuously with the same particle of fluid. This procedure, of course, is not feasible, but we do not need to make measurements with moving instruments because Eq. (1-11.2) gives the relation between the derivative dQ/dt and the derivative $\partial Q/\partial t$.

1-12 Differentiation of a Vector Field

If $\mathbf{Q}(x, y, z, t)$ is a vector field, Eq. (1-11.2) remains valid; that is,

$$\frac{d\mathbf{Q}}{dt} = u \frac{\partial \mathbf{Q}}{\partial x} + v \frac{\partial \mathbf{Q}}{\partial y} + w \frac{\partial \mathbf{Q}}{\partial z} + \frac{\partial \mathbf{Q}}{\partial t} \quad (1-12.1)$$

This follows from the fact that Eq. (1-11.2) is valid for each of the components of the vector $\mathbf{Q}$. Equation (1-12.1) may be written in the form

$$\frac{d\mathbf{Q}}{dt} = (\mathbf{q} \cdot \nabla)\mathbf{Q} + \frac{\partial \mathbf{Q}}{\partial t} \quad (1-12.2)$$

If $\mathbf{Q} = \mathbf{q}$, $d\mathbf{Q}/dt$ is the acceleration vector $\mathbf{a}$. Consequently,

$$\mathbf{a} = \frac{d\mathbf{q}}{dt} = u \frac{\partial \mathbf{q}}{\partial x} + v \frac{\partial \mathbf{q}}{\partial y} + w \frac{\partial \mathbf{q}}{\partial z} + \frac{\partial \mathbf{q}}{\partial t} \quad (1-12.3)$$

or

$$\mathbf{a} = (\mathbf{q} \cdot \nabla)\mathbf{q} + \frac{\partial \mathbf{q}}{\partial t} \quad (1-12.4)$$

Thus, the acceleration field is derived from the velocity field.

1-13 Curl of a Vector Field

Let $\mathbf{q} = iu + jv + kw$ be a vector field. Then $\nabla \times \mathbf{q}$ is a vector field that is denoted by curl $\mathbf{q}$. Hence, by Eq. (1-7.13),

$$\text{curl } \mathbf{q} = \nabla \times \mathbf{q} = \begin{vmatrix} \mathbf{i} & \mathbf{j} & \mathbf{k} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ u & v & w \end{vmatrix} \quad (1-13.1)$$

or

$$\text{curl } \mathbf{q} = \mathbf{i} \left(\frac{\partial w}{\partial y} - \frac{\partial v}{\partial z} \right) + \mathbf{j} \left(\frac{\partial u}{\partial z} - \frac{\partial w}{\partial x} \right) + \mathbf{k} \left(\frac{\partial v}{\partial x} - \frac{\partial u}{\partial y} \right) \quad (1-13.2)$$

It can be shown that the vector field curl $\mathbf{q}$ is independent of the choice of coordinates. A physical significance is later attributed to curl $\mathbf{q}$ if $\mathbf{q}$ denotes the velocity of a fluid. Curl $\mathbf{q}$ may also be related to the rotation of a volume element of a deformable body (see Chapter 2).

1-14 Eulerian Continuity Equation for Fluids

Let $\mathbf{q} = iu + jv + kw$ be an unsteady *velocity* field of a compressible fluid. Let us consider the rate of mass flow out of a space cell $dx dy dz = dV$ fixed with respect to (x, y, z) axes (see Fig. 1-14.1). The mass that flows in through the face AB during a time interval dt is $\rho u dy dz dt$, where ρ is the mass density. The mass that flows out through the face CD during dt is $\{\rho u + [\partial(\rho u)/\partial x] dx\} dy dz dt$. Similar expressions are obtained for the mass flows out of the other pairs of faces. Accordingly, the net mass that passes out of the cell dV during dt is

$$\left[\frac{\partial(\rho u)}{\partial x} + \frac{\partial(\rho v)}{\partial y} + \frac{\partial(\rho w)}{\partial z} \right] dV dt \quad (a)$$

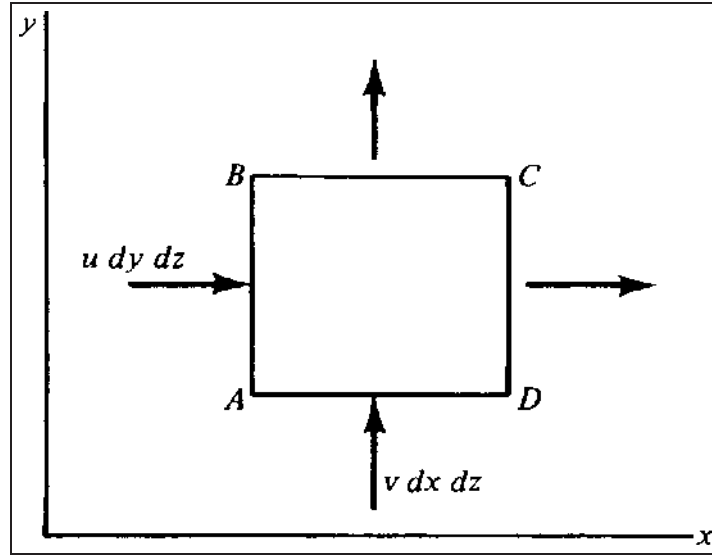


Figure 1-14.1

With the differential operator ∇ [see Eq. (1-8.5)] this may be written as

$$\nabla \cdot (\rho \mathbf{q}) dV dt \quad (b)$$

The product $\rho \mathbf{q}$ is called *current density*.

If $\mathbf{a}(x, y, z; t)$ is any vector field, $\nabla \cdot \mathbf{a}$ is called the *divergence* of the field. Accordingly, the notation $\text{div } \mathbf{a}$ is sometimes used to denote $\nabla \cdot \mathbf{a}$. Note that $\text{div } \mathbf{a}$ is a scalar. Accordingly, by Eq. (b), the mass that flows out of the volume element dV during dt is

$$dV dt \text{div}(\rho \mathbf{q}) \quad (c)$$

The name “divergence” originates in this physical idea.

Because mass is conserved in the velocity field of a fluid, the mass that passes into the fixed cell dV during time dt equals the increase of mass in the cell during dt . Now, the mass in the cell at the time t is ρdV . Consequently, the increase of mass during dt is

$$\frac{\partial \rho}{\partial t} dV dt \quad (d)$$

Because Eq. (d) must be the negative of Eq. (c), we obtain

$$\frac{\partial \rho}{\partial t} + \text{div}(\rho \mathbf{q}) = 0 \quad (1-14.1)$$

Equation (1-14.1) is known as the *Eulerian⁴ continuity equation* for fluids. Any real velocity field must conform to this relation. For steady flow, the term $\partial\rho/\partial t$ disappears.

For an incompressible fluid, $\rho = \text{constant}$. Consequently, the Eulerian form of the continuity equation for an incompressible fluid takes the simpler form:

$$\text{div } \mathbf{q} = 0 \quad \text{or} \quad \frac{\partial u}{\partial x} + \frac{\partial v}{\partial y} + \frac{\partial w}{\partial z} = 0 \quad (1-14.2)$$

This is valid even for unsteady flow of an incompressible fluid. Liquids may usually be considered to be incompressible except in the study of compression waves.

The case in which the velocity $\mathbf{q}$ is the gradient of a scalar function has great theoretical importance, that is, the case where

$$\mathbf{q} = -\text{grad } \phi \quad (1-14.3)$$

where $\phi(x, y, z; t)$ is a scalar function. The flow is then said to be *irrotational* or *derivable from a potential function* ϕ . Then the velocity component in the direction of a unit vector $\mathbf{n}$ is

$$q_n = \mathbf{q} \cdot \mathbf{n} = -\mathbf{n} \cdot \text{grad } \phi \quad (1-14.4)$$

Hence, by Eq. (1-8.3),

$$q_n = -\frac{d\phi}{ds} \quad (1-14.5)$$

That is, q_n is equal to the negative of the directional derivative of ϕ in the direction $\mathbf{n}$.

Equation (1-14.3) may be written

$$u = -\frac{\partial \phi}{\partial x} \quad v = -\frac{\partial \phi}{\partial y} \quad w = -\frac{\partial \phi}{\partial z}$$

Accordingly, by Eq. (1-14.2) the continuity equation for irrotational flow of an incompressible fluid is

$$\nabla^2 \phi = 0 \quad (1-14.6)$$

Thus, the continuity equation for irrotational flow of an incompressible fluid reduces to the Laplace equation (see Section 1-8). A general expression for the Laplace equation in orthogonal curvilinear coordinates in three-dimensional space is derived in Section 1-22.

⁴This form of the equation of continuity is referred to as the spatial form in modern continuum mechanics (see Chapter 2).

1-15 Divergence Theorem

Let $\mathbf{a}(x, y, z)$ be any continuous and differentiable vector field. We may regard $\mathbf{a}$ as current density in a hypothetical fluid. Then, by Eq. (c) of Section 1-14, $\text{div } \mathbf{a} \, dx \, dy \, dz$ is the net rate at which fluid flows out of the fixed space element $dx \, dy \, dz$. Hence, if R is a given fixed region of space that is bounded by a surface S , the net rate at which fluid passes out of R is

$$\iiint_R \text{div } \mathbf{a} \, dx \, dy \, dz$$

This must also be the rate at which fluid passes through the surface S . If dS is an element of area of this surface with outward-directed unit normal $\mathbf{n}$, the rate of flow through dS is $\mathbf{a} \cdot \mathbf{n} \, dS$. Hence,

$$\iiint_R \text{div } \mathbf{a} \, dx \, dy \, dz = \iint_S \mathbf{a} \cdot \mathbf{n} \, dS \quad (1-15.1)$$

Thus, a volume integral is transformed into a surface integral.

Equation (1-15.1) is known as the *divergence theorem* (also *Gauss's theorem*). It is purely mathematical; the reference to flow is simply an artifice to facilitate the derivation. Rigorous mathematical derivations of the theorem are given in books on advanced calculus (Goursat, 2005).

If (U, V, W) are the components of the vector $\mathbf{a}$, Eq. (1-15.1) may be expressed in scalar form:

$$\iiint_R \left(\frac{\partial U}{\partial x} + \frac{\partial V}{\partial y} + \frac{\partial W}{\partial z} \right) dx \, dy \, dz = \iint_S (Un_1 + Vn_2 + Wn_3) \, dS = \iint_S a_n \, dS \quad (1-15.2)$$

where a_n denotes the projection of $\mathbf{a}$ in the direction of $\mathbf{n}$, and (n_1, n_2, n_3) are the direction cosines of the unit vector $\mathbf{n}$; the functions (U, V, W) are unrestricted, aside from the requirements of continuity and differentiability. The surface S may consist of a finite number of smooth parts that are joined together along edges. If the vector $\mathbf{n}$ is directed inward, the sign of the right side of Eq. (1-15.2) is reversed.

Many useful results can be obtained by giving special forms to the functions (U, V, W) . For example, if $U = AB$, $V = W = 0$, we obtain

$$\iiint_R A \frac{\partial B}{\partial x} \, dx \, dy \, dz = - \iiint_R B \frac{\partial A}{\partial x} \, dx \, dy \, dz + \iint_S ABn_1 \, dS \quad (1-15.3)$$

Corresponding results for y and z are obtained by setting $V = AB$, $U = W = 0$, and so on. These equations are similar in form to the formula for integration by

parts of a single integral. Alternatively, if we take $V = W = 0$, Eq. (1-15.2) yields

$$\iiint_R \frac{\partial U}{\partial x} dx dy dz = \iint_R U n_1 dS \quad (1-15.3a)$$

Similar results are obtained for $U = W = 0$ and $U = V = 0$. Equation (1-15.3a) is called *Gauss's theorem*. More generally, *Gauss's theorem* may be written in the form

$$\int_V \frac{\partial F_i}{\partial x_i} dV = \int_S F_i n_i dS \quad i = 1, 2, 3 \quad (1-15.3b)$$

where $F_i = F_i(x_1, x_2, x_3)$, V denotes volume, S denotes surface of volume V with unit normal vector $\mathbf{n} : (n_1, n_2, n_3)$, and $x_1 \equiv x$, $x_2 \equiv y$, and $x_3 \equiv z$.

Another useful relation may be obtained as follows: Let $\mathbf{a}$ be the product of a scalar ϕ and a vector $\mathbf{A}$; that is,

$$\mathbf{a} = \phi \mathbf{A}$$

Then

$$\text{div } \mathbf{a} = \phi \text{ div } \mathbf{A} + \frac{\partial \phi}{\partial x} A_x + \frac{\partial \phi}{\partial y} A_y + \frac{\partial \phi}{\partial z} A_z \quad (1-15.4)$$

or

$$\text{div } \mathbf{a} = \phi \text{ div } \mathbf{A} + (\text{grad } \phi) \cdot \mathbf{A}$$

Accordingly, Eq. (1-15.1) yields

$$\iint_S \phi A_n dS = \iiint_R [\phi \text{ div } \mathbf{A} + (\text{grad } \phi) \cdot \mathbf{A}] dV \quad (1-15.5)$$

If, furthermore, the vector $\mathbf{A}$ is representable as the gradient of a scalar function ψ ($\mathbf{A} = \text{grad } \psi$), then by Eq. (1-14.5), $A_n = d\psi/dn$ and

$$\text{div } \mathbf{A} = \frac{\partial^2 \psi}{\partial x^2} + \frac{\partial^2 \psi}{\partial y^2} + \frac{\partial^2 \psi}{\partial z^2} = \nabla^2 \psi$$

Hence, for $\mathbf{A} = \text{grad } \psi$, Eq. (1-15.5) becomes

$$\iint_S \phi \frac{\partial \psi}{\partial n} dS = \iiint_R [\phi \nabla^2 \psi + (\text{grad } \phi) \cdot (\text{grad } \psi)] dV \quad (1-15.6)$$

Equation (1-15.6) holds for any two functions ϕ and ψ that are finite, continuous, and twice differentiable in R .

If we subtract from Eq. (1-15.6) the equation obtained by interchanging ϕ and ψ , we obtain

$$\iint_S \left(\phi \frac{\partial \psi}{\partial n} - \psi \frac{\partial \phi}{\partial n} \right) dS = \iiint_R (\phi \nabla^2 \psi - \psi \nabla^2 \phi) dV \quad (1-15.7)$$

Both Eqs. (1-15.6) and (1-15.7) are referred to as *Green's theorem*. They find extensive use in mathematical physics.

The above results are useful in transformations from volume to surface integrals and vice versa.

1-16 Divergence Theorem in Two Dimensions

The two-dimensional analog of Eq. (1-15.2) is

$$\iint_R \left(\frac{\partial U}{\partial x} + \frac{\partial V}{\partial y} \right) dx dy = \oint_C (U n_1 + V n_2) ds \quad (1-16.1)$$

where U and V are any continuous and differentiable functions of (x, y) . Here R denotes a region of the (x, y) plane, and C is the curve that bounds the region R (Fig. 1-16.1). The unit normal vector (n_1, n_2) is directed outward. The element of arc length of the curve C is denoted by ds . The circle on the integral sign shows that the integration extends completely around the curve C , in the counterclockwise sense.

Referring to the figure, we have $n_1 = \cos \alpha$, and $n_2 = \sin \alpha$. Hence, $n_1 ds = dy$, and $n_2 ds = -dx$, where (dx, dy) is the displacement along the curve C , corresponding to the increment ds . Hence, by Eq. (1-16.1),

$$\iint_R \left(\frac{\partial U}{\partial x} + \frac{\partial V}{\partial y} \right) dx dy = \oint_C (U dy - V dx) \quad (1-16.2)$$

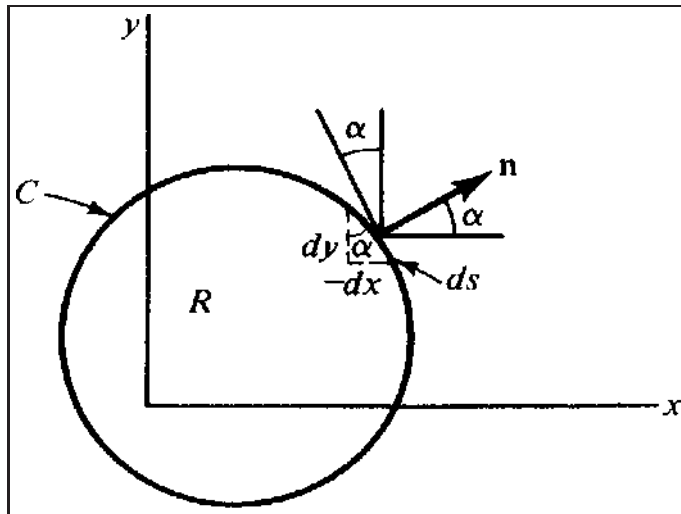


Figure 1-16.1

This relation is sometimes called *Green's theorem of the plane*. Another form of Green's theorem is obtained by the substitution $U = v$, $V = -u$. Then

$$\iint_R \left(\frac{\partial v}{\partial x} - \frac{\partial u}{\partial y} \right) dx dy = \oint_C (u dx + v dy) \quad (1-16.3)$$

With $U = AB$, $V = 0$, Eq. (1-16.2) yields

$$\iint_R A \frac{\partial B}{\partial x} dx dy = - \iint_R B \frac{\partial A}{\partial x} dx dy + \oint_C AB dy \quad (1-16.4)$$

Furthermore, analogous to the three-dimensional development of Eqs. (1-15.6) and (1-15.7), we have

$$\oint_C \phi \frac{\partial \psi}{\partial n} ds = \iint_R [\phi \nabla^2 \psi + (\text{grad } \phi) \cdot (\text{grad } \psi)] dx dy \quad (1-16.5)$$

$$\oint_C \left(\phi \frac{\partial \psi}{\partial n} - \psi \frac{\partial \phi}{\partial n} \right) ds = \iint_R (\phi \nabla^2 \psi - \psi \nabla^2 \phi) dx dy \quad (1-16.6)$$

where (ϕ, ψ) are functions of (x, y) only.

1-17 Line and Surface Integrals (Application of Scalar Product)

Line Integral. Consider a vector $\mathbf{F}$ defined at each point on a curve C (Fig. 1-17.1). The vector $\mathbf{F}$ forms an angle α with the tangent to the curve C at point P . In general, the vector $\mathbf{F}$ may vary in magnitude and direction along the curve. Let s be an arc length measured along the curve. The length of an infinitesimal element of the curve at point P is ds . The vector $d\mathbf{s}$ with magnitude ds is directed along the tangent line to the curve at point P (Fig. 1-17.1).

By Eq. (1-7.8), the projection of the vector $\mathbf{F}$ along the tangent to the curve is $\mathbf{F} \cdot d\mathbf{s} = F(\cos \alpha) ds$. The integral

$$\int_C \mathbf{F} \cdot d\mathbf{s} = \int_C F(\cos \alpha) ds \quad (1-17.1)$$

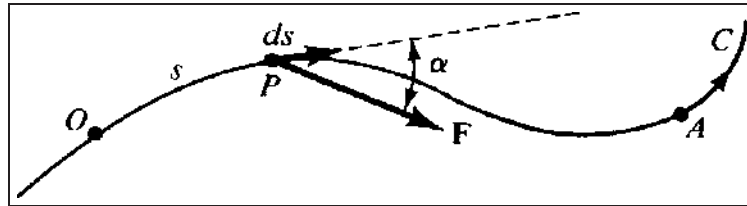


Figure 1-17.1

is called the line integral of the vector $\mathbf{F}$ along the curve C . The C in Eq. (1-17.1) denotes integration along the curve C . By Eq. (1-17.1) it is apparent that the line integral of a vector is the integral of the tangential component of the vector taken along a path.

The line integral Eq. (1-17.1) finds numerous applications in physical problems. For example, if $\mathbf{F}$ denotes a force that acts on a particle P that travels along curve C , the line integral of the tangential component of $\mathbf{F}$ from point O to point A represents the work performed by the force $\mathbf{F}$ as the particle travels from O to A . If $\mathbf{F}$ denotes the electric field intensity, that is, the force that acts on a unit charge in an electric field, the line integral between any two points represents the potential difference between the two points. If $\mathbf{F}$ denotes the velocity at any point in a fluid, the line integral taken around a closed path in the fluid represents the *circulation* of the fluid.

Surface Integral. In Section 1-15 it was shown that the volume of fluid that passes through a surface S in a unit time is

$$\iint_S \mathbf{q} \cdot \mathbf{n} dS = \iint_S q_n dS \quad (1-17.2)$$

where $\mathbf{q}$ is the velocity field and $\mathbf{n}$ is the unit normal to the surface. This integral is called the *surface integral* of the vector $\mathbf{q}$. Accordingly, the expression *surface integral of a vector* denotes the integral of the normal component of the vector over a surface.

1-18 Stokes's Theorem

Equation (1-16.3) may be written as

$$\oint_C \mathbf{q} \cdot d\mathbf{r} = \iint_R \mathbf{n} \cdot \text{curl } \mathbf{q} dS \quad (1-18.1)$$

where $d\mathbf{r} = (dx, dy)$, $\mathbf{q}$ denotes the vector (u, v, w) , and $\mathbf{n}$ now denotes the unit normal to the plane area R [directed in the positive z direction, if the coordinates (x, y, z) are right handed]. Although Eq. (1-18.1) has been proven only if R is a region in the (x, y) plane, it remains valid if R is any plane area in space with any orientation, for Eq. (1-18.1) is invariant under a coordinate transformation; that is, Eq. (1-18.1) does not depend on the choice of coordinates.

Our result may be generalized still further. The curve C need not be a plane curve; it may be any closed space curve, and R may be any surface S that caps this curve. Any capping surface of the curve C may be divided into infinitesimal cells. Each cell is a plane element of area. Consequently, Eq. (1-18.1) applies for any one of the cells. We may then sum Eq. (1-18.1) over all cells. Then the right side of the equation simply becomes the surface integral of $\text{curl } \mathbf{q}$ over the entire capping surface S of curve C . On the left side we have the sum of line integrals of

$\mathbf{q}$ about the boundaries of the cells. However, the line integrals over the boundaries of contiguous cells cancel, as any inner boundary of a cell is described twice, only in the positive sense and once in the negative sense. Consequently, only the line integral on the outer boundary C remains.

Accordingly, we have Stokes's theorem: *The line integral of a vector field about any closed curve equals the surface integral of the normal component of the curl of the vector over any capping surface.*

If $\mathbf{q}$ is a velocity field, then $\text{curl } \mathbf{q}$ is called the *vorticity vector*. Consequently, in the terminology of fluid mechanics Stokes's theorem is expressed as follows: *The circulation on any closed curve equals the flux of vorticity through the loop.*

1-19 Exact Differential

Let $M(x, y)$ and $N(x, y)$ be two functions of x and y such that $M, N, \partial M/\partial y$, and $\partial N/\partial x$ are continuous and single valued at every point of a simply connected⁵ region. The differential expression $M dx + N dy$ is said to be *exact* if there exists a function $f(x, y)$ such that $df = M dx + N dy$. Now, by definition,

$$df = \frac{\partial f}{\partial x} dx + \frac{\partial f}{\partial y} dy \quad (1-19.1)$$

Consequently, if $M dx + N dy$ is exact, $M = \partial f/\partial x$, and $N = \partial f/\partial y$. Therefore,

$$\frac{\partial M}{\partial y} = \frac{\partial N}{\partial x} \quad \text{or} \quad \frac{\partial N}{\partial x} - \frac{\partial M}{\partial y} = 0 \quad (1-19.2)$$

Accordingly, Eq. (1-19.2) is a *necessary condition* for $M dx + N dy$ to be an exact differential.

Equation (1-19.2) is also a *sufficient condition*. Assume that Eq. (1-19.2) is satisfied. Set

$$F(x, y) = \int M dx$$

where integration is performed with respect to x . Then $\partial F/\partial x = M$ and

$$\frac{\partial^2 F}{\partial x \partial y} = \frac{\partial M}{\partial y} = \frac{\partial N}{\partial x}$$

Therefore,

$$\frac{\partial}{\partial x} \left(N - \frac{\partial F}{\partial y} \right) = 0 \quad \text{or} \quad N = \frac{\partial F}{\partial y} + g(y)$$

Set $f(x, y) = F(x, y) + \int g(y) dy$. Then $N = \partial f/\partial y$ and $M = \partial F/\partial x = \partial f/\partial x$. Hence, $M dx + N dy = df$; that is, $M dx + N dy$ is an exact differential.

⁵A simply connected region has the property that any closed curve drawn on it can, by continuous deformation, be shrunk to a point without crossing the boundary of the region. For the significance of simple connectivity, see Courant (1992), Vol. II.

If $f = f(x, y, z)$, $df = P(x, y, z)dx + Q(x, y, z)dy + R(x, y, z)dz$ where $P = \partial f / \partial x$, $Q = \partial f / \partial y$, and $R = \partial f / \partial z$, an argument analogous to the two-dimensional case leads to the necessary and sufficient conditions that df be an exact differential in the form

$$\frac{\partial Q}{\partial x} - \frac{\partial P}{\partial y} = 0 \quad \frac{\partial R}{\partial y} - \frac{\partial Q}{\partial z} = 0 \quad \frac{\partial P}{\partial z} - \frac{\partial R}{\partial x} = 0 \quad (1-19.3)$$

1-20 Orthogonal Curvilinear Coordinates in Three-Dimensional Space

Let three independent scalar functions (u, v, w) be defined in terms of three independent variables (x, y, z) as follows:

$$u = U(x, y, z) \quad v = V(x, y, z) \quad w = W(x, y, z) \quad (1-20.1)$$

By independent functions, we mean that Eqs. (1-20.1) yield unique solutions for (x, y, z) :

$$x = X(u, v, w), \quad y = Y(u, v, w), \quad z = Z(u, v, w) \quad (1-20.2)$$

For example, if (x, y, z) represents rectangular Cartesian coordinates, and (u, v, w) represents cylindrical coordinates, Eq. (1-20.2) is of the form

$$x = u \cos v \quad y = u \sin v \quad z = w \quad (1-20.3)$$

If (u, v, w) represents spherical coordinates, Eq. (1-20.2) is of the form

$$x = u \sin v \cos w \quad y = u \sin v \sin w \quad z = u \cos v \quad (1-20.4)$$

If (u, v, w) are assigned constant values, Eq. (1-20.1) becomes

$$\begin{aligned} U_0(x, y, z) &= \text{const} = u_0 \\ V_0(x, y, z) &= \text{const} = v_0 \\ W_0(x, y, z) &= \text{const} = w_0 \end{aligned} \quad (1-20.5)$$

Equations (1-20.5) represent three surfaces in space, called *coordinate surfaces*. The intersection of any two of these surfaces (say, $U_0 = u_0$ and $V_0 = v_0$) determines a curve in space, the *w curvilinear coordinate line*. The *u* and *v* curvilinear coordinate lines are defined similarly. The three surfaces $U_0 = u_0$, $V_0 = v_0$, and $W_0 = w_0$ intersect at a point in space. Hence, a point in space is associated with each triplet (u_i, v_i, w_i) .

If the three systems of surfaces defined by triplets (u_i, v_i, w_i) are mutually perpendicular (i.e., if the curvilinear coordinate lines through any point are mutually perpendicular), the curvilinear coordinate system is said to be *orthogonal*.

A very special case of an orthogonal curvilinear coordinate system is the rectangular Cartesian coordinate system. For rectangular coordinates,

$$x = u \quad y = v \quad z = w$$

Hence, three coordinate surfaces are the mutually perpendicular planes:

$$x = u_0 \quad y = v_0 \quad z = w_0$$

The intersection of any two of these planes is a *coordinate line*; for example, the intersection of planes $x = u_0$, $y = v_0$ determines a z coordinate line. Cylindrical coordinates [Eq. (1-20.3)] and spherical coordinate [Eq. (1-20.4)] are also examples of orthogonal curvilinear coordinate systems. Another example is elliptic coordinates.

1-21 Expression for Differential Length in Orthogonal Curvilinear Coordinates

Let ($\mathbf{i}, \mathbf{j}, \mathbf{k}$) be unit vectors along (x, y, z) axes, respectively. Let (u, v, w) be a system of orthogonal curvilinear coordinates. Let v and w be constant. Then at any point the tangent vector to the u coordinate line is

$$\mathbf{U} = x_u \mathbf{i} + y_u \mathbf{j} + z_u \mathbf{k} \quad (1-21.1)$$

where the u subscript denotes partial differentiation. Similarly, tangent vectors to the v and w coordinate lines are

$$\mathbf{V} = x_v \mathbf{i} + y_v \mathbf{j} + z_v \mathbf{k} \quad \mathbf{W} = x_w \mathbf{i} + y_w \mathbf{j} + z_w \mathbf{k} \quad (1-21.2)$$

Vectors $\mathbf{U}$, $\mathbf{V}$, $\mathbf{W}$ are mutually perpendicular. Hence, by the scalar product definition of two vectors,

$$\mathbf{U} \cdot \mathbf{V} = \mathbf{V} \cdot \mathbf{W} = \mathbf{W} \cdot \mathbf{U} = 0 \quad (1-21.3)$$

Also, if (h_1, h_2, h_3) are the magnitudes of the lengths of vectors ($\mathbf{U}, \mathbf{V}, \mathbf{W}$), respectively, the scalar product definition yields

$$h_1^2 = \mathbf{U} \cdot \mathbf{U} \quad h_2^2 = \mathbf{V} \cdot \mathbf{V} \quad h_3^2 = \mathbf{W} \cdot \mathbf{W} \quad (1-21.4)$$

Hence, by Eqs. (1-20.2), (1-21.1), (1-21.2), and (1-21.4), $h_1 = h_1(u, v, w)$, $h_2 = h_2(u, v, w)$, and $h_3 = h_3(u, v, w)$.

Consider a line element PQ , where $P = P(x, y, z)$ and $Q = Q(x + dx, y + dy, z + dz)$. The differential length ds of the line element PQ is given by the relation

$$ds^2 = dx^2 + dy^2 + dz^2 \quad (1-21.5)$$

By Eq. (1-20.2),

$$\begin{aligned} dx &= x_u du + x_v dv + x_w dw \\ dy &= y_u du + y_v dv + y_w dw \\ dz &= z_u du + z_v dv + z_w dw \end{aligned} \quad (1-21.6)$$

Substituting Eqs. (1-21.6) into Eq. (1-21.5) and utilizing Eqs. (1-21.2), (1-21.3), and (1-21.4), we obtain

$$ds^2 = h_1^2 du^2 + h_2^2 dv^2 + h_3^2 dw^2 \quad (1-21.7)$$

Equation (1-21.7) expresses the differential length ds in terms of the orthogonal curvilinear coordinates (u, v, w) . The coefficients (h_1, h_2, h_3) are called *Lamé coefficients*. The Lamé coefficients are equal in magnitude to the lengths of the vectors $(\mathbf{U}, \mathbf{V}, \mathbf{W})$ tangent to (u, v, w) coordinate lines, respectively. The quantities (h_1^2, h_2^2, h_3^2) are known as the *components of the metric tensor of space* (Synge and Schild, 1978).

1-22 Gradient and Laplacian in Orthogonal Curvilinear Coordinates

Consider the infinitesimal parallelepiped whose diagonal is the line element ds . The faces of the parallelepiped coincide with the planes $u = \text{constant}$, $v = \text{constant}$, $w = \text{constant}$ (Fig. 1-22.1).

The gradient u , (∇u) has the direction normal to the surface $u = \text{constant}$; that is, the direction of $\mathbf{U}$ or the direction of the unit vector $\mathbf{U}/|\mathbf{U}| = \mathbf{U}/h_1$ [see

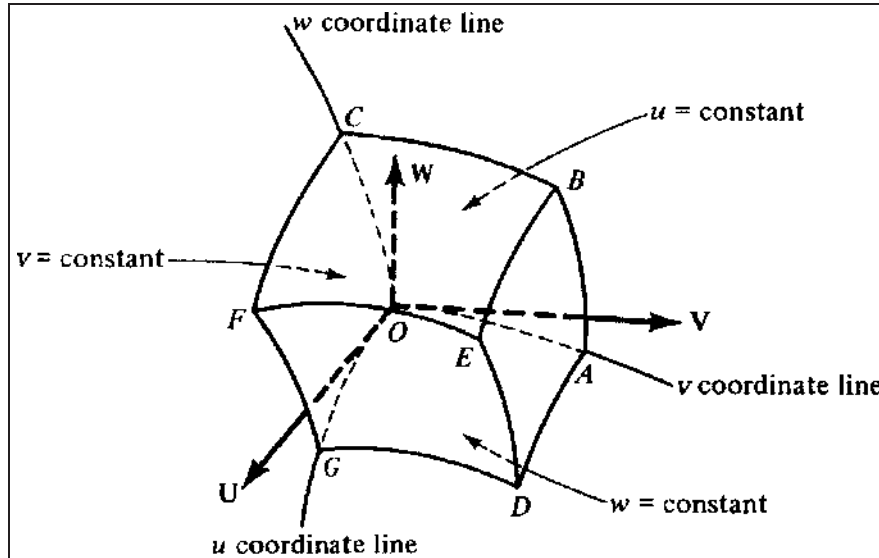


Figure 1-22.1

Eq. (1-21.4)]. The magnitude of ∇u is equal to the derivative of u in this direction. Hence, by Eq. (1-21.7), with v and w constant, the magnitude of ∇u is

$$\frac{du}{ds} = \frac{1}{h_1} \quad (1-22.1)$$

Hence, the gradient vector is

$$\nabla u = \frac{1}{h_1^2} \mathbf{U} \quad (1-22.2)$$

Similarly, the gradient of v and w are

$$\nabla v = \frac{1}{h_2^2} \mathbf{V} \quad \nabla w = \frac{1}{h_3^2} \mathbf{W} \quad (1-22.3)$$

By the definition of ∇ and by the rule for partial differentiation, that is,

$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial u} \frac{\partial u}{\partial x} + \frac{\partial f}{\partial v} \frac{\partial v}{\partial x} + \frac{\partial f}{\partial w} \frac{\partial w}{\partial x}$$

if $f(u, v, w)$ is any scalar point function, then the gradient of f is

$$\nabla f = \frac{\partial f}{\partial u} \nabla u + \frac{\partial f}{\partial v} \nabla v + \frac{\partial f}{\partial w} \nabla w \quad (1-22.4)$$

Substituting Eqs. (1-22.2) and (1-22.3) into Eq. (1-22.4), we obtain

$$\nabla f = \frac{1}{h_1} \frac{\partial f}{\partial u} \mathbf{u} + \frac{1}{h_2} \frac{\partial f}{\partial v} \mathbf{v} + \frac{1}{h_3} \frac{\partial f}{\partial w} \mathbf{w} \quad (1-22.5)$$

where $(\mathbf{u}, \mathbf{v}, \mathbf{w})$ are unit vectors in the directions of $(\mathbf{U}, \mathbf{V}, \mathbf{W})$, respectively; that is,

$$\mathbf{u} = \frac{\mathbf{U}}{h_1} \quad \mathbf{v} = \frac{\mathbf{V}}{h_2} \quad \mathbf{w} = \frac{\mathbf{W}}{h_3} \quad (1-22.6)$$

Equation (1-22.5) represents the gradient of a scalar in orthogonal curvilinear coordinates. Consequently, by Eq. (1-22.5), the expression for the operator ∇ in orthogonal curvilinear coordinates is

$$\nabla = \frac{1}{h_1} \mathbf{u} \frac{\partial}{\partial u} + \frac{1}{h_2} \mathbf{v} \frac{\partial}{\partial v} + \frac{1}{h_3} \mathbf{w} \frac{\partial}{\partial w} \quad (1-22.7)$$

To derive the expression for the Laplacian ∇^2 , we first derive the expression for the divergence of a vector field, $\mathbf{Q} = (Q_1, Q_2, Q_3)$, that is, $\nabla \cdot \mathbf{Q}$, in orthogonal curvilinear coordinates.

Consider again the infinitesimal parallelepiped of Fig. 1-22.1. The lengths of its edges are $h_1 du$, $h_2 dv$, and $h_3 dw$, and its volume is $h_1 h_2 h_3 du dv dw$. To facilitate

the calculation of the divergence of $\mathbf{Q}$, we use Green's theorem for transforming volume integrals into surface integrals:

$$\iiint_{\text{through volume}} (\nabla \cdot \mathbf{Q}) dV = \iint_{\text{over bounding surface}} \mathbf{Q} \cdot \mathbf{n} dS \quad (1-22.8)$$

The contribution of the surface $OABC$, Fig. (1-22.1), to the integral over the surface of the parallelepiped taken in the direction of the outward normal is $-Q_1 h_2 dv h_3 dw$. The contribution of the surface $DEFG$ is

$$Q_1 h_2 h_3 dv dw + \frac{\partial}{\partial u}(Q_1 h_2 h_3) du dv dw$$

Hence, the net contribution of the coordinate surfaces perpendicular to u coordinate lines is

$$\frac{\partial}{\partial u}(Q_1 h_2 h_3) du dv dw \quad (1-22.9)$$

Similarly, the contributions of the coordinate surfaces perpendicular to v and w coordinate lines, respectively, are

$$\frac{\partial}{\partial v}(Q_2 h_1 h_3) du dv dw \quad \frac{\partial}{\partial w}(Q_3 h_1 h_2) du dv dw \quad (1-22.10)$$

Because the volume of the infinitesimal parallelepiped, Fig. 1-22.1, is infinitesimal,

$$\lim_{v \rightarrow 0} \iiint (\nabla \cdot \mathbf{Q}) dV \rightarrow \nabla \cdot \mathbf{Q} h_1 h_2 h_3 du dv dw \quad (1-22.11)$$

Consequently, by Eqs. (1-22.8) to (1-22.11),

$$\nabla \cdot \mathbf{Q} = \frac{1}{h_1 h_2 h_3} \left[\frac{\partial}{\partial u}(Q_1 h_2 h_3) + \frac{\partial}{\partial v}(Q_2 h_1 h_3) + \frac{\partial}{\partial w}(Q_3 h_1 h_2) \right] = \text{div } \mathbf{Q} \quad (1-22.12)$$

Equation (1-22.12) represents the formula for the divergence of a vector field $\mathbf{Q}$ in terms of general three-dimensional orthogonal curvilinear coordinates.

Setting $\nabla f = \mathbf{Q}$ and noting by Eq. (1-22.5) that $Q_1 = (1/h_1)(\partial f/\partial u)$, $Q_2 = (1/h_2)(\partial f/\partial v)$, and $Q_3 = (1/h_3)(\partial f/\partial w)$, we obtain, by Eqs. (1-22.6) and (1-22.12),

$$\nabla^2 f = \nabla \cdot \nabla f = \frac{1}{h_1 h_2 h_3} \left[\frac{\partial}{\partial u} \left(\frac{h_2 h_3}{h_1} \frac{\partial f}{\partial u} \right) + \frac{\partial}{\partial v} \left(\frac{h_1 h_3}{h_2} \frac{\partial f}{\partial v} \right) + \frac{\partial}{\partial w} \left(\frac{h_1 h_2}{h_3} \frac{\partial f}{\partial w} \right) \right] \quad (1-22.13)$$

Equation (1-22.13) represents the Laplacian of a scalar function $f(u, v, w)$ in general three-dimensional orthogonal curvilinear coordinates. Hence, the Laplace equation $\nabla^2 f = 0$ in general three-dimensional orthogonal curvilinear coordinates is obtained by setting the right-hand side of Eq. (1-22.13) equal to zero.

For plane (two-dimensional) orthogonal curvilinear coordinates, $h_3 = 1$ and $\partial/\partial w = 0$.

PART III ELEMENTS OF TENSOR ALGEBRA

1-23 Index Notation: Summation Convention

Gibbs vector notation may be considered to replace and extend conventional scalar notation. For example, the scalar representation (F_x, F_y, F_z) of a force with respect to rectangular Cartesian axes is fully replaced by the vector notation $\mathbf{F}$. Likewise, index notation may be considered to replace and extend Gibbs vector notation. Thus, the vector $\mathbf{F}$ may be represented by the symbol F_i , where the subscript (index) i is understood to take values 1, 2, 3 (or the values x, y, z). Hence, the notation F_i is equivalent to (F_1, F_2, F_3) or to (F_x, F_y, F_z) , where subscripts (1, 2, 3) or subscripts (x, y, z) denote projections of the force along rectangular Cartesian coordinate axes (1, 2, 3) or (x, y, z) .

Restricting ourselves to rectangular Cartesian coordinates, we indicate coordinates by indices (1, 2, 3) instead of letters (x, y, z) . For example, the coordinate of a general point X in (x, y, z) space are denoted by $x_i = (x_1, x_2, x_3)$ or more briefly by x_i , with the understanding that i takes the values (1, 2, 3). The coordinates of a specific point P are denoted by p_i , the letter p identifying the point and the index i , the separate coordinates (see Fig. 1-23.1). Similarly, axes (x, y, z) may be denoted by (x_1, x_2, x_3) , or simply by x_i . Axes x_i may also be denoted by the notations (01, 02, 03) or (1, 2, 3).

The direction cosines of a line L with respect to axes x_i are denoted by $\alpha_1, \alpha_2, \alpha_3$ or briefly by α_i . Any other letter may replace α . For example, the direction cosines of line L may also be denoted by β_i , by m_i , by n_i , and so on.

The sum of two vectors q_i, r_i is $q_i + r_i$. The scalar product of two vectors u_α, v_α is [see Eq. (1-7.6)]

$$\mathbf{u} \cdot \mathbf{v} = u_1 v_1 + u_2 v_2 + u_3 v_3 = \sum_{\alpha=1}^3 u_\alpha v_\alpha \quad (1-23.1)$$

Equation (1-23.1) may be simplified by the use of conventional *summation* notation. For example, we may write Eq. (1-23.1) in the form

$$\mathbf{u} \cdot \mathbf{v} = u_\alpha v_\alpha \quad (1-23.2)$$

with the understanding that the *repeated Greek index* α implies summation over the values (1, 2, 3). Accordingly, if m_α and n_α denote the direction cosines of two

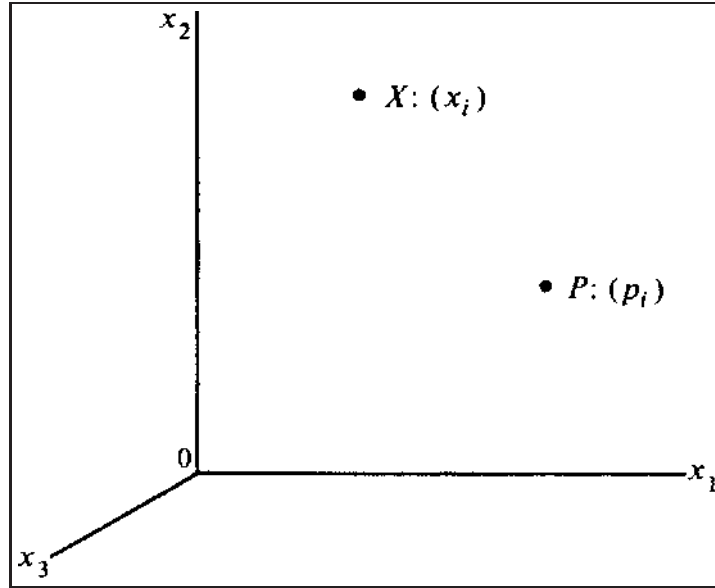


Figure 1-23.1

unit vectors directed along two lines M and N in (x, y, z) space, by the scalar product of vectors, the angle θ between lines M and N is given by the relation [see Eq. (1-7.8) and the discussion following it]

$$\cos \theta = m_\alpha n_\alpha \quad (1-23.3)$$

If lines M and N coincide, $\theta = 0$. Then Eq. (1-23.3) yields (with $m_\alpha = n_\alpha$)

$$m_1^2 + m_2^2 + m_3^2 = 1 \quad (1-23.4)$$

Accordingly, the sum of the squares of the direction cosines of a directed line in (x, y, z) space is equal to 1 [see Eq. (1-7.9)].

In general, a repeated index that is to be summed will be denoted by a Greek letter. We thus avoid the necessity of using some special notation for a repeated index that is not summed. Because the operation of summing is independent of the Greek index used to denote the summation process, the following representations of $\cos \theta$ are equivalent [see Eq. (1-23.3)]:

$$\cos \theta = m_\alpha n_\alpha = m_\beta n_\beta = m_\gamma n_\gamma = \cdots$$

as each of the representations denotes $m_1 n_1 + m_2 n_2 + m_3 n_3$. Accordingly, a repeated Greek index is called a *summing index* or a *dummy index*. An index that appears only once in a general term is called a *free index*. Thus, in the term

$A_{\alpha\beta\beta}$, the index β is a dummy index and the index α is a free index, the value of α being independent of the values of β . For example, if we assign the value 1 to α , the term $A_{\alpha\beta\beta}$ represents the sum $A_{111} + A_{122} + A_{133}$.

If a *repeated* index is *not* to be summed, we denote it by a *Latin letter* ($a, b, c, \dots, z$). Thus, $m_i n_i$ denotes any element of the set $(m_1 n_1, m_2 n_2, m_3 n_3)$, depending on the values assigned to i . For example, if $i = 2$, then $m_i n_i$ denotes the element $m_2 n_2$.

If several dummy indexes occur in a general term, summation is implied for each index separately. For example,

$$\begin{aligned} x_{i\alpha\beta} y_{\alpha\beta} &= x_{i1\beta} y_{1\beta} + x_{i2\beta} y_{2\beta} + x_{i3\beta} y_{3\beta} \\ &= x_{i11} y_{11} + x_{i12} y_{12} + x_{i13} y_{13} \\ &\quad + x_{i21} y_{21} + x_{i22} y_{22} + x_{i23} y_{23} \\ &\quad + x_{i31} y_{31} + x_{i32} y_{32} + x_{i33} y_{33} \end{aligned}$$

Thus, for every value of the free index i , there are nine terms in the sum $x_{i\alpha\beta} y_{\alpha\beta}$.

In modern algebra, the range of the index is often extended from $(1, 2, 3)$ to $(1, 2, 3, \dots, n)$. Thus, we may write

$$A_{i\alpha} x_{\alpha} = A_{i1} x_1 + A_{i2} x_2 + \dots + A_{in} x_n$$

where the summing index α takes values $(1, 2, 3, \dots, n)$.

To avoid confusion, an index already appearing in a general term as a free index should not be used as a dummy index, as no meaning is given indexes that appear more than twice. Thus, notations such as $A_{\beta\beta} x_{\beta}$ should be avoided. For example, if $x = A_{\alpha} y_{\alpha}$ and $y_i = B_{i\alpha} z_{\alpha}$, the expression for x in terms of (z_1, z_2, z_3) is written

$$x = A_{\alpha} B_{\alpha\beta} z_{\beta}$$

not in the meaningless form

$$x = A_{\alpha} B_{\alpha\alpha} z_{\alpha}$$

Rectangular Arrays. A set of numbers arranged in the following form is called a *rectangular array*:

$$\begin{bmatrix} a_{11} & a_{12} & a_{13} & \cdots & a_{1n} \\ a_{21} & a_{22} & a_{23} & \cdots & a_{2n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & a_{m3} & \cdots & a_{mn} \end{bmatrix} \quad (1-23.5)$$

where, in general, $m \neq n$.

More generally, such an array of numbers is called a *matrix*. In the study of matrix theory, extensive rules are laid down for the multiplication of matrices (Section 1-28). However, the role of products in matrix theory is to a large extent replaced by summation convention. A typical element of an array is denoted

by a_{ij} , the index i referring to the i th row of the array and the index j to the j th column. For brevity, the entire array [Eq. (1-23.5)] is denoted by

$$[a_{ij}] \quad (1-23.6)$$

If $m = n$, the array is called a *square array*. In the theory of continuous media, we are concerned primarily with square arrays.

If the arrays $[a_{ij}]$, $[b_{ij}]$, $[c_{ij}]$, $\dots$ all have the same number of rows and the same number of columns, a linear combination $[h_{ij}]$, of $[a_{ij}]$, $[b_{ij}]$, $[c_{ij}]$, $\dots$ is defined by the elements

$$h_{ij} = Aa_{ij} + Bb_{ij} + Cc_{ij} + \dots \quad (1-23.7)$$

where $A, B, C, \dots$ are arbitrary constants independent of i and j . In particular, the *sum* $[a_{ij} + b_{ij} + c_{ij}]$ of the three arrays $[a_{ij}]$, $[b_{ij}]$, and $[c_{ij}]$ has the typical element $a_{ij} + b_{ij} + c_{ij}$.

A square array $[a_{ij}]$ is said to be *symmetric* if

$$a_{ij} = a_{ji} \quad (1-23.8)$$

for all pairs of values of i, j ; a square array is said to be *skew symmetric* or *antisymmetric* if

$$a_{ij} = -a_{ji} \quad (1-23.9)$$

for all pairs of i, j . For an antisymmetric array, it follows, by Eq. (1-23.9), that $a_{ii} = a_{jj} = 0$.

An *arbitrary square array* (neither symmetric nor antisymmetric) may be represented as the sum of a symmetric array and an antisymmetric array. For example, any two numbers r and s can always be written in the form

$$r = \frac{1}{2}(x + y) \quad s = \frac{1}{2}(x - y)$$

by letting

$$x = r + s \quad y = r - s$$

Hence, we may express a typical element of the arbitrary square array $[a_{ij}]$ in the form

$$\begin{aligned} a_{ij} &= \frac{1}{2}(a_{ij} + a_{ji}) + \frac{1}{2}(a_{ji} - a_{ji}) \\ &= \frac{1}{2}(a_{ij} + a_{ji}) + \frac{1}{2}(a_{ij} - a_{ji}) \end{aligned}$$

or

$$a_{ij} = c_{ij} + d_{ij} \quad (1-23.10)$$

where

$$c_{ij} = \frac{1}{2}(a_{ij} + a_{ji}) = c_{ji}$$

denotes the elements of a symmetric square array, and

$$d_{ij} = \frac{1}{2}(a_{ij} - a_{ji}) = -d_{ji}$$

denotes the elements of an antisymmetric square array.

1-24 Transformation of Tensors under Rotation of Rectangular Cartesian Coordinate System

In this section we consider briefly some tensor transformations and properties that are important in the theory of deformable media. For simplicity, we restrict our discussion to rectangular Cartesian coordinates. Accordingly, the results presented here are special cases of more general tensor transformations (Synge and Schild, 1978; Spain, 2003).

Let (x, y, z) and (X, Y, Z) denote two right-handed rectangular Cartesian coordinate systems with common origin (Fig. 1-24.1). The cosines of the angles between the six coordinate axes may be represented in tabular form (Table 1-24.1). Each entry in Table 1-24.1 is the cosine of the angle between the two coordinate axes

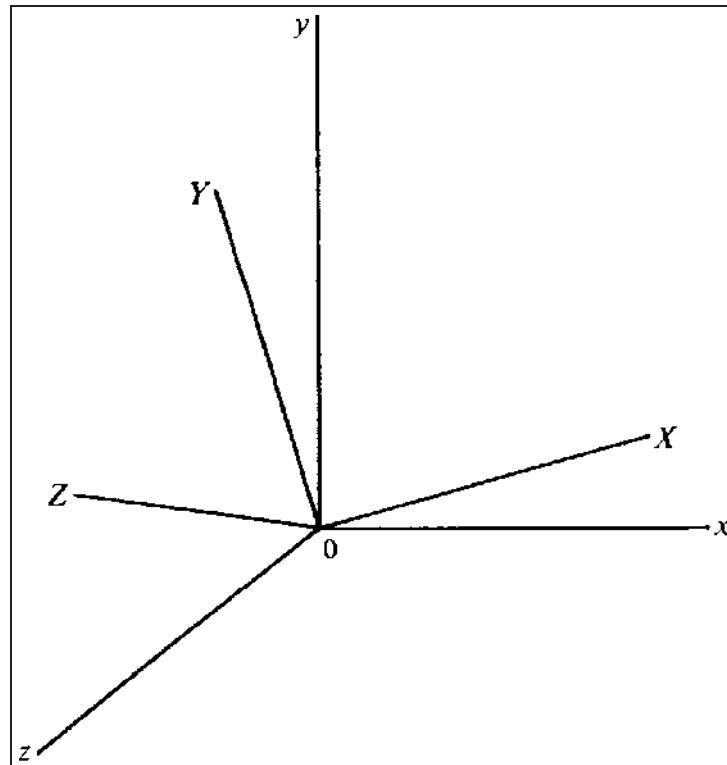


Figure 1-24.1

TABLE 1-24.1

	x	y	z
X	a_{11}	a_{12}	a_{13}
Y	a_{21}	a_{22}	a_{23}
Z	a_{31}	a_{32}	a_{33}

designated at the top of its column and left of its row. For example, a_{23} denotes the cosine of the angle between the Y axis and the z axis; that is, $a_{\alpha\beta}$ represents the direction cosines of the angle between the axes designated by the row α and the column β of Table 1-24.1. Because the elements of Table 1-24.1 are direction cosines, they satisfy the following relations (Eisenhart, 2005):

$$\begin{aligned} a_{1\beta}^2 + a_{2\beta}^2 + a_{3\beta}^2 &= 1 & \beta &= 1, 2, 3 \\ a_{\alpha 1}^2 + a_{\alpha 2}^2 + a_{\alpha 3}^2 &= 1 & \alpha &= 1, 2, 3 \end{aligned} \quad (1-24.1)$$

Equation (1-24.1) signifies that the sum of the squares of the elements of any row or column of Table 1-24.1 is 1. Furthermore, because the axes (X, Y, Z) are mutually perpendicular, we have

$$a_{\alpha 1}a_{\beta 1} + a_{\alpha 2}a_{\beta 2} + a_{\alpha 3}a_{\beta 3} = 0 \quad \alpha, \beta = 1, 2, 3 \quad \alpha \neq \beta \quad (1-24.2)$$

Similarly, because (x, y, z) are mutually perpendicular, we have further

$$a_{1\beta}a_{1\alpha} + a_{2\beta}a_{2\alpha} + a_{3\beta}a_{3\alpha} = 0 \quad \alpha, \beta = 1, 2, 3 \quad \alpha \neq \beta \quad (1-24.3)$$

Equations (1-24.2) and (1-24.3) signify that the sum of the products of corresponding elements in any two rows or any two columns in Table 1-24.1 is zero. In other words, they express the orthogonality of axes (X, Y, Z) and the orthogonality of axes (x, y, z). For this reason, they are called *orthogonality relations*.

Another important relation between the coefficients of Table 1-24.1 may be obtained as follows. Noting that the direction cosines of a unit vector with respect to (x, y, z) axes are identical to the projections of the unit vector on the coordinate axes, we regard the direction cosines (a_{11}, a_{12}, a_{13}) as the components on (x, y, z) axes of a unit vector in the X direction. Similarly, (a_{21}, a_{22}, a_{23}) and (a_{31}, a_{32}, a_{33}) represent unit vectors in the Y direction and the Z direction, respectively. Hence, by the vector product of vectors [see Eq. (1-7.13)], if the two coordinate systems (x, y, z) and (X, Y, Z) are both right handed (or both left handed), we obtain the vector relation

$$(a_{11}, a_{12}, a_{13}) = (a_{21}, a_{22}, a_{23}) \times (a_{31}, a_{32}, a_{33})$$

or, in scalar notation,

$$\begin{aligned} a_{11} &= a_{22}a_{33} - a_{23}a_{32} \\ a_{12} &= a_{31}a_{23} - a_{21}a_{33} \\ a_{13} &= a_{21}a_{32} - a_{22}a_{31} \end{aligned} \quad (1-24.4)$$

Similar relations hold for $(a_{21}, a_{22}, a_{23}), \dots, (a_{13}, a_{23}, a_{33})$. In index notation, the entire set of relations may be written

$$a_{kr} = a_{ip}a_{jq} - a_{iq}a_{jp} \quad (1-24.5)$$

where (i, j, k) , the first indexes of each direction cosine, may take any cyclic order of 1, 2, 3, 1, 2, $\dots$, and where (p, q, r) , the second indexes of each direction cosine, take independently any cyclic order of 1, 2, 3, 1, 2, $\dots$. For example, let (i, j, k) be (2, 3, 1) and let (p, q, r) be (2, 3, 1). Then Eq. (1-24.5) yields

$$a_{11} = a_{22}a_{33} - a_{23}a_{32}$$

Similarly, $(i, j, k) = (1, 2, 3), (p, q, r) = (3, 1, 2)$ yields

$$a_{32} = a_{13}a_{21} - a_{11}a_{23}$$

Equations (1-24.5) are also referred to as orthogonality relations, as they express the orthogonality of axes (x, y, z) and of axes (X, Y, Z) .

In view of Eqs. (1-24.4), the second equation of Eqs. (1-24.1), with $\alpha = 1$, may be written

$$a_{11}(a_{22}a_{33} - a_{23}a_{32}) + a_{12}(a_{31}a_{23} - a_{21}a_{33}) + a_{13}(a_{21}a_{32} - a_{22}a_{31}) = 1$$

Similar expressions hold for $\alpha = 2, 3; \beta = 1, 2, 3$.

In determinant notation, the above equation may be written in the form

$$\det a_{\alpha\beta} = \begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} = 1 \quad (1-24.6)$$

where $\det$ denotes determinant. If the coordinate system is left handed, it may be shown that $\det a_{\alpha\beta} = -1$. Consequently, we have the following theorem:

Theorem 1-24.1. *Any one of the direction cosines of a set of right-handed (left-handed) rectangular Cartesian axes measured with respect to a second set of right-handed (left-handed) rectangular Cartesian axes is equal to its cofactor (the negative of its cofactor) in the determinant formed from the square array of direction cosines [see Eqs. (1-24.4) and (1-24.6)]. Furthermore, the numerical value of the determinant is 1 (−1).*

In the following, we consider right-handed coordinate systems only.

Let the coordinates of a point P be (x, y, z) with respect to axes (x, y, z) . Then, with respect to (X, Y, Z) axes, the coordinates of P may be expressed in terms of coordinates (x, y, z) by the equations

$$\begin{aligned} X &= a_{11}x + a_{12}y + a_{13}z \\ Y &= a_{21}x + a_{22}y + a_{23}z \\ Z &= a_{31}x + a_{32}y + a_{33}z \end{aligned} \quad (1-24.7)$$

For (X, Y, Z) axes with origin at (a_{10}, a_{20}, a_{30}) , Eqs. (1-24.7) may be generalized by the substitution $X = X - a_{10}$, $Y = Y - a_{20}$, and $Z = Z - a_{30}$.

Conversely, with respect to (x, y, z) axes, the coordinates of P expressed in terms of (X, Y, Z) are given by the relations (because $\det a_{\alpha\beta} = 1$)

$$\begin{aligned} x &= a_{11}X + a_{21}Y + a_{31}Z \\ y &= a_{12}X + a_{22}Y + a_{32}Z \\ z &= a_{13}X + a_{23}Y + a_{33}Z \end{aligned} \quad (1-24.8)$$

With the summation notation introduced in Section 1-23, Eq. (1-24.7) becomes

$$X_\alpha = a_{\alpha 1}x_1 + a_{\alpha 2}x_2 + a_{\alpha 3}x_3 \quad \alpha = 1, 2, 3 \quad (1-24.9)$$

or

$$X_\alpha = a_{\alpha\beta}x_\beta \quad \alpha, \beta = 1, 2, 3$$

Similarly, Eq. (1-24.8) may be written

$$\begin{aligned} x_\beta &= a_{1\beta}X_1 + a_{2\beta}X_2 + a_{3\beta}X_3 \quad \beta = 1, 2, 3 \\ x_\beta &= a_{\alpha\beta}X_\alpha \quad \alpha, \beta = 1, 2, 3 \end{aligned} \quad (1-24.10)$$

For given values of α and β , the value of $a_{\alpha\beta}$ in Eq. (1-24.9) is identical to the value of $a_{\alpha\beta}$ in Eq. (1-24.10). This follows from the definition of the entries in Table 1-24.1.

With the understanding that α, β take values 1, 2, 3, Eqs. (1-24.9) and (1-24.10) are written

$$X_\alpha = a_{\alpha\beta}x_\beta \quad (1-24.11)$$

and

$$x_\beta = a_{\alpha\beta}X_\alpha \quad (1-24.12)$$

Because a repeated Greek index is always summed, it may be replaced by any convenient letter, as noted in Section 1-23. Accordingly, the following forms for Eq. (1-24.11) are all equivalent:

$$X_\alpha = a_{\alpha\beta}x_\beta = a_{\alpha\gamma}x_\gamma = a_{\alpha\zeta}x_\zeta$$

Scalars. Quantities such as temperature and density that may be represented by a single number—for example, 10°C or 30 g/cm^3 —are called scalars. Under a transformation of coordinate axes, scalars remain unchanged; that is, scalars are invariant under coordinate transformations. For this reason, scalars are often called invariants. In tensor theory, scalars are called *tensors of zero order*.

Vectors. In summation notation a vector is represented by the symbol u_i (Section 1-23). Suppose the arrow OP representing the vector u_i is attached

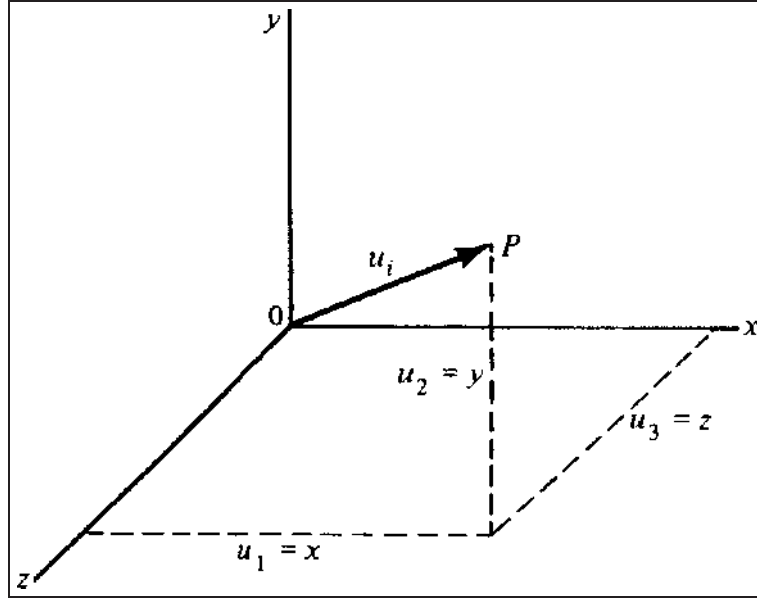


Figure 1-24.2

to a rectangular Cartesian coordinate system (x, y, z) , as in Fig. 1-24.2. Then the coordinates of P correspond to the components (u_1, u_2, u_3) of vector u_i . Consequently, under a transformation from one rectangular Cartesian coordinate system to another, the components of a three-dimensional vector transform according to the relationship [see Eq. (1-24.11)]

$$U_\alpha = a_{\alpha\beta} u_\beta \quad (1-24.13)$$

The vector u_i remains fixed in space. Such sets of three components (i.e., vectors) are called *tensors of first order*. Tensors of first order require only one index for their representation. Multiplication of a first-order tensor by a zero-order tensor (i.e., multiplication of a vector by a scalar) yields another first-order tensor. For example, multiplication of u_i by a constant c yields cu_i . Hence, by Eq. (1-24.13), $a_{\alpha\beta}(cu_\beta) = c(a_{\alpha\beta}u_\beta) = cU_\alpha$. Thus, cu_i is a tensor of first order, as it obeys the rules of transformation of a tensor of first order. Furthermore, the addition of two tensors of first order (two vectors) yields a tensor of first order (a vector). For example, if u_p, v_p are two tensors of first order, by Eq. (1-24.13) we have

$$U_\alpha = a_{\alpha\beta} u_\beta \quad V_\alpha = a_{\alpha\beta} v_\beta$$

Addition of these equations yields

$$U_\alpha + V_\alpha = a_{\alpha\beta} u_\beta + a_{\alpha\beta} v_\beta = a_{\alpha\beta} (u_\beta + v_\beta)$$

Hence, $u_\beta + v_\beta$ is a tensor of first order, as it transforms according to Eq. (1-24.13).

Tensors of Higher Order. Multiplication of tensors of first order leads to quantities that are not tensors of zero or first order. For example, let u_ζ and v_η be two first-order tensors in the rectangular Cartesian coordinate system (x, y, z) . Let U_α , V_β denote the corresponding tensors in the rectangular Cartesian coordinate system (X, Y, Z) . Then, by Eq. (1-24.13),

$$U_\alpha V_\beta = (a_{\alpha\zeta} u_\zeta)(a_{\beta\eta} v_\eta) = a_{\alpha\zeta} a_{\beta\eta} u_\zeta v_\eta \quad (1-24.14)$$

or

$$W_{\alpha\beta} = a_{\alpha\zeta} a_{\beta\eta} w_{\zeta\eta}$$

where $W_{\alpha\beta} = U_\alpha V_\beta$ and $w_{\zeta\eta} = u_\zeta v_\eta$ represent the products of the vectors U_α , V_β in the (X, Y, Z) system and u_ζ , v_η in the (x, y, z) system, respectively.

Because both ζ and η are dummy indexes, for given values of α , β the right-hand side of Eq. (1-24.14) contains nine terms. Accordingly, Eq. (1-24.14) represents nine equations, each with nine terms. Quantities that transform according to Eq. (1-24.14) are called *tensors of second order*. In the symbolical representation of tensors of second order, two indexes are required. Many quantities other than the product of two vectors transform according to Eq. (1-24.14). For example, components of stress and of strain transform according to Eq. (1-24.14) under a change of rectangular coordinate systems (see Chapters 2 and 3). Accordingly, the components of stress and of strain form second-order tensors.

In a similar fashion, a *tensor of third order* is formed by multiplying together three first-order tensors, and so on. Thus, an n th-order tensor may be formed by multiplying together n first-order tensors. Essentially, this means that we have available means of specifying components of n th-order tensors with respect to any set of rectangular Cartesian axes and rules for transforming these components to any other set of rectangular Cartesian axes. Hence, the statement that a quantity is a tensor quantity may be proved by comparison with these known tensor transformations. For example, this technique was employed in the proof that the sum of two first-order tensors yields a first-order tensor.

In summary, a tensor of zero order (scalar) is a single quantity that depends on position in space but not on the coordinate system. A tensor of first order (vector) is a quantity whose components transform according to Eq. (1-24.13). Hence, with respect to a rectangular Cartesian coordinate system in three-dimensional space, a tensor of first order contains $3^1 = 3$ elements of components. A tensor of second order is a quantity that transforms according to Eq. (1-24.14). With respect to rectangular Cartesian coordinate systems in three-dimensional space, a second-order tensor has $3^2 = 9$ elements.

A tensor of n th order is a quantity whose components transform according to the rule⁶

$$T_{p_1 p_2 \dots p_n} = a_{p_1 q_1} a_{p_2 q_2} \dots a_{p_n q_n} t_{q_1 q_2 \dots q_n} \quad (1-24.15)$$

⁶See Synge and Schild (1978). Here we let dummy indexes be denoted by $q_1, q_2, \dots, q_n$.

With respect to rectangular Cartesian coordinate axes in three-dimensional space, an n th-order tensor has 3^n elements. Thus, a fourth-order tensor has 81 elements, a fifth-order tensor has 243 elements, and a tenth-order tensor has 59,049 elements.

In general developments of continuous-media mechanics, fourth-order tensors play a prominent role (Green and Zerna, 2002).

1-25 Symmetric and Antisymmetric Parts of a Tensor

If we interchange α and β in Eq. (1-24.14), we obtain

$$W_{\beta\alpha} = a_{\beta\zeta} a_{\alpha\eta} w_{\zeta\eta} \quad (1-25.1)$$

Because ζ and η are dummy indexes, we may interchange them. Thus, Eq. (1-25.1) may be written

$$W_{\beta\alpha} = a_{\beta\eta} a_{\alpha\zeta} w_{\eta\zeta} = a_{\alpha\zeta} a_{\beta\eta} w_{\eta\zeta} \quad (1-25.2)$$

Hence, comparing Eqs. (1-24.14) and (1-25.2), we see that $w_{\eta\zeta}$ transforms according to the same rule as $w_{\zeta\eta}$. The tensor $w_{\zeta\eta}$ is said to be *conjugate to* $w_{\eta\zeta}$. Thus, if $w_{\zeta\eta}$ is a tensor of second order, another tensor of second order is obtained by interchanging η and ζ . Consequently, $w_{\zeta\eta} + w_{\eta\zeta}$ and $w_{\zeta\eta} - w_{\eta\zeta}$ are tensors of second order. Symbolically, we may represent the tensors $w_{\zeta\eta}$ and $w_{\eta\zeta}$ as follows:

$$w_{\zeta\eta} = \begin{pmatrix} w_{11} & w_{12} & w_{13} \\ w_{21} & w_{22} & w_{23} \\ w_{31} & w_{32} & w_{33} \end{pmatrix}$$

and

$$w_{\eta\zeta} = \begin{pmatrix} w_{11} & w_{21} & w_{31} \\ w_{12} & w_{22} & w_{32} \\ w_{13} & w_{23} & w_{33} \end{pmatrix}$$

Then

$$\begin{aligned} w_{\zeta\eta} + w_{\eta\zeta} &= \begin{pmatrix} 2w_{11} & w_{12} + w_{21} & w_{13} + w_{31} \\ w_{21} + w_{12} & 2w_{22} & w_{23} + w_{32} \\ w_{31} + w_{13} & w_{32} + w_{23} & 2w_{33} \end{pmatrix} \\ &= w_{\eta\zeta} + w_{\zeta\eta} \end{aligned} \quad (1-25.3)$$

and

$$\begin{aligned} w_{\zeta\eta} - w_{\eta\zeta} &= \begin{pmatrix} 0 & w_{12} - w_{21} & w_{13} - w_{31} \\ w_{21} - w_{12} & 0 & w_{23} - w_{32} \\ w_{31} - w_{13} & w_{32} - w_{23} & 0 \end{pmatrix} \\ &= -(w_{\eta\zeta} - w_{\zeta\eta}) \end{aligned} \quad (1-25.4)$$

Because $w_{\zeta\eta} + w_{\eta\zeta}$ is unaltered by interchanging ζ and η , it is called a symmetrical tensor of second order. However, when ζ and η are interchanged in $w_{\zeta\eta} - w_{\eta\zeta}$,

each element changes in sign. Hence, $w_{\zeta\eta} - w_{\eta\zeta}$ is called an antisymmetrical tensor of second order. Also, by Eqs. (1-25.3) and (1-25.4),

$$w_{\zeta\eta} = \frac{1}{2}(w_{\zeta\eta} + w_{\eta\zeta}) + \frac{1}{2}(w_{\zeta\eta} - w_{\eta\zeta}) = S_{\zeta\eta} + A_{\zeta\eta} \quad (1-25.5)$$

where $S_{\zeta\eta}$ is a symmetric second-order tensor and $A_{\zeta\eta}$ is antisymmetric. Consequently, a second-order tensor may be resolved into symmetric and antisymmetric parts. Furthermore, because the antisymmetric part contains only three components, $w_{12} - w_{21}$, $w_{13} - w_{31}$, $w_{23} - w_{32}$, it may be associated with a vector u_i . Equation (1-25.5) is analogous to Eq. (1-23.10).

Problem. Let $w_{\zeta\eta} + w_{\eta\zeta} = 2S_{\zeta\eta} = 2S_{\eta\zeta}$ and $w_{\zeta\eta} - w_{\eta\zeta} = A_{\zeta\eta} = -(w_{\eta\zeta} - w_{\zeta\eta}) = -A_{\eta\zeta}$, where $w_{\zeta\eta}$ is a tensor of second order. Show that the product of the symmetric tensor $S_{\zeta\eta}$ and the antisymmetric tensor $A_{\zeta\eta}$ vanishes; that is, show that $S_{\zeta\eta}A_{\zeta\eta} = 0$.

1-26 Symbols δ_{ij} and ϵ_{ijk} (the Kronecker Delta and the Alternating Tensor)

The use of the following notation often simplifies the writing of equations:

$$\delta_{ij} = \begin{cases} 1 & \text{for } i = j \\ 0 & \text{for } i \neq j \end{cases} \quad (1-26.1)$$

The symbol δ_{ij} is called the *Kronecker delta*.

Using the notation δ_{ij} with respect to axes (x, y, z) , we may write the second of Eqs. (1-24.1) and Eqs. (1-24.2) collectively as

$$a_{\alpha\gamma}a_{\beta\gamma} = \delta_{\alpha\beta} \quad (1-26.2)$$

Similarly, with respect to axes (X, Y, Z) we may express the first of Eqs. (1-24.1) and Eqs. (1-24.3) in the form

$$a_{\gamma\beta}a_{\gamma\alpha} = \delta_{\beta\alpha} \quad (1-26.3)$$

The Kronecker delta has the following important properties:

1. $\delta_{\lambda\lambda} = \delta_{11} + \delta_{22} + \delta_{33} = 3$
2. $\delta_{i\lambda}\delta_{j\lambda} = \delta_{ij}$
3. $p_{i\lambda}\delta_{j\lambda} = p_{ij}$

Property 3 is a generalization of 2. It is called the *rule of substitution of indexes*, as the multiplication of $\delta_{j\lambda}$ substitutes the index j for the index λ .

The set of quantities δ_{ij} , $i, j = 1, 2, 3$ constitutes a *tensor of the second order*. To prove this, we must show that δ_{ij} transforms according to Eq. (1-24.14) under a

transformation of rectangular Cartesian axes. The array δ_{ij} consists of the elements $\delta_{11} = 1$, $\delta_{22} = 1$, $\delta_{33} = 1$, $\delta_{12} = 0$, $\delta_{23} = 0$, and $\delta_{13} = 0$. Accordingly, if we set $\delta_{\sigma\gamma} = w_{\sigma\gamma}$ and substitute in Eq. (1-24.14), we get

$$W_{\alpha\beta} = \delta'_{\alpha\beta} = a_{\alpha\sigma}a_{\beta\gamma}\delta_{\sigma\gamma} = a_{\alpha 1}a_{\beta 1} + a_{\alpha 2}a_{\beta 2} + a_{\alpha 3}a_{\beta 3}$$

Hence, by Eqs. (1-26.1) and (1-26.2),

$$\delta'_{\alpha\beta} = \begin{cases} 1 & \text{for } \alpha = \beta \\ 0 & \text{for } \alpha \neq \beta \end{cases}$$

Thus, it follows that the array ($\delta_{11} = \delta_{22} = \delta_{33} = 1$, $\delta_{12} = \delta_{13} = \delta_{23} = 0$) is transformed into itself by the tensor transformation Eq. (1-24.14). This transformation is in accord with the definition of Eq. (1-26.1). Hence, $\delta_{\alpha\beta}$ is a second-order tensor. A tensor whose respective components (elements) are the same with respect to all sets of coordinate systems is called an *isotropic tensor*. In view of the fact that δ_{ij} is a tensor and in view of the substitution property 3 above, δ_{ij} is sometimes referred to as the *substitution tensor*.

Symbol ϵ_{ijk} . The symbol ϵ_{ijk} is defined as follows:

$$\epsilon_{ijk} \begin{cases} 1 & \text{if } i, j, k \text{ are in cyclic order } 1, 2, 3, 1, 2, \dots \\ 0 & \text{if any two of } i, j, k \text{ are equal} \\ -1 & \text{if } i, j, k \text{ are in anticyclic order } 3, 2, 1, 2, 3, \dots \end{cases} \quad (1-26.4)$$

For example,

$$\begin{aligned} \epsilon_{123} &= \epsilon_{312} = \epsilon_{231} = 1 \\ \epsilon_{112} &= \epsilon_{121} = \epsilon_{322} = \dots = 0 \\ \epsilon_{321} &= \epsilon_{213} = \epsilon_{132} = -1 \end{aligned} \quad (1-26.5)$$

By definition of δ_{ij} and ϵ_{ijk} , it follows that

$$\epsilon_{ijk}\delta_{ij} = \epsilon_{iik} = 0 \quad \text{no summation} \quad (1-26.6)$$

Furthermore, it follows by Eqs. (1-26.5) and (1-26.6) that

$$\epsilon_{\alpha\beta k}\delta_{\alpha\beta} = 0 \quad \text{summed} \quad (1-26.7)$$

In terms of ϵ_{ijk} , the orthogonality relations [Eq. (1-24.5)] may be written

$$\epsilon_{ij\alpha}a_{\alpha n} = \epsilon_{\alpha\beta n}a_{i\alpha}a_{j\beta} \quad (1-26.8)$$

where i, j, n take independently any value 1, 2, 3. The proof of Eq. (1-26.8) is left for the problems.

The array ϵ_{ijk} transforms according to the rules of transformation of a third-order isotropic tensor. To show this, we note that a third-order tensor transforms

according to the rule [see Eq. (1-24.15)]

$$T_{ijk} = a_{i\alpha}a_{j\beta}a_{k\gamma}t_{\alpha\beta\gamma} \quad (1-26.9)$$

Hence, we must show that $\epsilon_{\alpha\beta\gamma}$ transforms according to the rule

$$\epsilon_{ijk} = a_{i\alpha}a_{j\beta}a_{k\gamma}\epsilon_{\alpha\beta\gamma} \quad (1-26.10)$$

Substituting Eq. (1-26.8) into the right side of Eq. (1-26.10), we obtain

$$a_{k\gamma}a_{\alpha\gamma}\epsilon_{ij\alpha}$$

But $a_{k\gamma}a_{\alpha\gamma} = \delta_{k\alpha}$, by Eq. (1-26.2). Hence,

$$a_{k\gamma}a_{\alpha\gamma}\epsilon_{ij\alpha} = \delta_{k\alpha}\epsilon_{ij\alpha} = \epsilon_{ijk}$$

Accordingly, Eq. (1-26.10) is verified. In view of the properties noted in Eqs. (1-26.4) and (1-26.10), the symbol ϵ_{ijk} is called the *alternating tensor*.

1-27 Homogeneous Quadratic Forms

The most general homogeneous quadratic form in the variables X_i , $i = 1, 2, 3$, may be written in index notation as

$$Q = a_{\alpha\beta}X_\alpha X_\beta \quad \alpha, \beta = 1, 2, 3 \quad (1-27.1)$$

where a_{ij} denotes the following square array of real elements:

$$\begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix} \quad (1-27.2)$$

The quadratic form Q written in expanded form is

$$\begin{aligned} Q = & a_{11}X_1^2 + a_{22}X_2^2 + a_{33}X_3^2 + (a_{12} + a_{21})X_1X_2 \\ & + (a_{13} + a_{31})X_1X_3 + (a_{23} + a_{32})X_2X_3 \end{aligned} \quad (1-27.3)$$

The determinant $\det a_{ij}$ is called the *determinant* of the array [Eq. (1-27.2)]. The expression (1-27.1) [or Eq. (1-27.3)] is called the *quadratic form* associated with the array $[a_{ij}]$. Without loss of generality, the array may be assumed symmetrical; that is, we may set $a_{ij} = a_{ji}$. Then Eq. (1-27.3) becomes

$$Q = a_{11}X_1^2 + a_{22}X_2^2 + a_{33}X_3^2 + 2a_{12}X_1X_2 + 2a_{13}X_1X_3 + 2a_{23}X_2X_3 \quad (1-27.4)$$

where we have simply replaced the notation $(a_{12} + a_{21})$ in Eq. (1-27.3) by $2a_{12}$ in Eq. (1-27.4), and so on.

The equation

$$\begin{vmatrix} a_{11} - r & a_{12} & a_{13} \\ a_{21} & a_{22} - r & a_{23} \\ a_{31} & a_{32} & a_{33} - r \end{vmatrix} = 0$$

or, in the index notation,

$$|a_{ij} - r\delta_{ij}| = 0 \quad (1-27.5)$$

is called the *characteristic* equation of the array $[a_{ij}]$. The three roots (r_1, r_2, r_3) of Eq. (1-27.5) are called the *characteristic roots*, or *latent roots*, or *eigenvalues* of the array $[a_{ij}]$ (Eisenhart, 2005; Hildebrand, 1992). In general, the r_i are distinct. However, special cases may occur in which two or all of the r_i are equal.

A necessary and sufficient condition that a set of linear algebraic equations

$$c_{i\alpha} X_\alpha = 0 \quad i = 1, 2, 3 \quad (1-27.6)$$

possess a solution other than the trivial solution $X_1 = X_2 = X_3 = 0$ is that the determinant of the coefficients $c_{i\alpha}$ of Eq. (1-27.6) vanishes (Pipes, 1959; Hildebrand, 1992). Accordingly,

$$|c_{i\alpha}| = 0 \quad (1-27.7)$$

represents a necessary and sufficient condition that Eq. (1-27.6) possess a solution $X_i (X_i \neq 0)$. Accordingly, by Eqs. (1-27.5), (1-27.6), and (1-27.7), it follows that for every r such that $|a_{ij} - r\delta_{ij}| = 0$, an array (X_i) exists such that

$$(a_{i\alpha} - r\delta_{i\alpha})X_\alpha = 0$$

Rewriting, we have

$$a_{i\alpha} X_\alpha = r\delta_{i\alpha} X_\alpha = r X_i \quad (1-27.8)$$

In other words, Eq. (1-27.5) expresses the necessary and sufficient condition that Eq. (1-27.8) possesses nontrivial solutions of X_i . The nontrivial solutions of Eq. (1-27.8) are called the *eigenvectors* of the array $[a_{ij}]$.

Let y_i denote any arbitrary array (y_1, y_2, y_3) . Then, by Eq. (1-27.8), we obtain the *bilinear* form

$$a_{\alpha\beta} X_\beta y_\alpha = r X_\alpha y_\alpha \quad (1-27.9)$$

If $y_i = X_i$, we obtain the *quadratic* form (Hildebrand, 1992)

$$a_{\alpha\beta} X_\alpha X_\beta = r X_\alpha X_\alpha \quad (1-27.10)$$

Orthogonality of Eigenvectors. Consider the case where the array X_i corresponds to the array m_i of direction cosines [Eq. (1-23.4)]. Assume that there exist two nonequal characteristic roots $r^{(1)}, r^{(2)}$ of Eq. (1-27.5). Then the corresponding solutions (eigenvectors) of Eq. (1-27.8) may be denoted by $m_i^{(1)}, m_i^{(2)}$.

Accordingly, Eq. (1-27.8) becomes

$$\begin{aligned} a_{i\alpha} m_{\alpha}^{(1)} &= r^{(1)} m_i^{(1)} & \text{for } r = r^{(1)} \\ a_{i\alpha} m_{\alpha}^{(2)} &= r^{(2)} m_i^{(2)} & \text{for } r = r^{(2)} \end{aligned} \quad (1-27.11)$$

Multiplying the first of Eqs. (1-27.11) by $m_i^{(2)}$ and the second by $m_i^{(1)}$ and subtracting, we obtain (because $a_{ij} = a_{ji}$)

$$[r^{(2)} - r^{(1)}] m_{\beta}^{(1)} m_{\beta}^{(2)} = 0$$

However, because by hypothesis $r^{(2)} \neq r^{(1)}$, it follows that

$$m_{\beta}^{(1)} m_{\beta}^{(2)} = 0 \quad (1-27.12)$$

Accordingly, the directions (eigenvectors) $m_{\beta}^{(1)}$, $m_{\beta}^{(2)}$ that correspond to the characteristic roots $r^{(1)}$ and $r^{(2)}$ are orthogonal. Furthermore, if $r^{(1)}$ and $r^{(2)}$ are two distinct characteristic roots and $m_{\beta}^{(1)}$ and $m_{\beta}^{(2)}$ are the corresponding direction cosines, by Eq. (1-27.10) we have

$$a_{\alpha\beta} m_{\alpha}^{(1)} m_{\beta}^{(2)} = r^{(2)} m_{\beta}^{(1)} m_{\beta}^{(2)} = r^{(1)} m_{\beta}^{(1)} m_{\beta}^{(2)} = 0$$

Hence

$$a_{\alpha\beta} m_{\alpha}^{(1)} m_{\beta}^{(2)} = 0 \quad (1-27.13)$$

This result is equivalent to the vanishing of shearing stress (or strain) components relative to principal axes (see Chapters 2 and 3).

Finally, note that the characteristic roots $r^{(1)}$, $r^{(2)}$ are real. We prove this by contradiction as follows: Assume that $r^{(1)}$ is complex. Denote its complex conjugate by $\bar{r}^{(1)}$. Then, taking the complex conjugate of the first of Eqs. (1-27.11), we obtain

$$a_{\alpha\beta} \bar{m}_{\beta}^{(1)} = \bar{r}^{(1)} \bar{m}_{\alpha}^{(1)} \quad (1-27.14)$$

Multiplying (1-27.14) by $m_{\alpha}^{(1)}$, we get

$$a_{\alpha\beta} m_{\alpha}^{(1)} \bar{m}_{\beta}^{(1)} = \bar{r}^{(1)} m_{\alpha}^{(1)} \bar{m}_{\alpha}^{(1)} \quad (1-27.15)$$

Multiplying the first of Eqs. (1-27.11) by $\bar{m}_{\alpha}^{(1)}$, we get

$$a_{\alpha\beta} \bar{m}_{\alpha}^{(1)} m_{\beta}^{(1)} = r^{(1)} m_{\alpha}^{(1)} \bar{m}_{\alpha}^{(1)} \quad (1-27.16)$$

Comparison of Eqs. (1-27.15) and (1-27.16) yields

$$[\bar{r}^{(1)} - r^{(1)}] m_{\alpha}^{(1)} \bar{m}_{\alpha}^{(1)} = 0$$

Because $m_\alpha^{(1)}\overline{m}_\alpha^{(1)}$ is the sum of squares of real numbers, it cannot be zero unless $m_1 = m_2 = m_3 = 0$. However, this is not possible because, by Eq. (1-23.4),

$$m_1^2 + m_2^2 + m_3^2 = 1$$

Hence

$$r^{(1)} = \overline{r}^{(1)}$$

That is, $r^{(1)}$ is equal to its conjugate $\overline{r}^{(1)}$. Accordingly, $r^{(1)}$ must be real.

1-28 Elementary Matrix Algebra

The matrix algebra outlined in this section plays an important role in modern structural analysis and in numerical methods of continuum mechanics such as finite element methods.

In Section 1-23 we noted that the rectangular array of m rows and n columns of numbers a_{ij} is called a matrix (more explicitly, an m by n matrix or a matrix of order m by n). The elements a_{ij} may be real or complex numbers or, more generally, may be matrices themselves. However, unless we state otherwise, we take the numbers a_{ij} to be real. In Section 1-23 we denoted the array by $[a_{ij}]$ and considered several properties of the array in terms of the individual elements a_{mn} . However, it is frequently more economical to treat a matrix as a single entity, particularly in algebraic operations involving addition, subtraction, multiplication, and division of several arrays. Accordingly, we employ the notation

$$A = [a_{ij}] \quad 1 \leq i \leq m \quad 1 \leq j \leq n \quad (1-28.1)$$

where A denotes the m by n matrix [Eq. (1-23.5)] of the m by n elements a_{ij} .

If $m = 1$,

$$A = [a_{11}, a_{12}, \dots, a_{1n}] = (a_{11}, a_{12}, \dots, a_{1n}) \quad (1-28.2)$$

contains one row. Hence it is called a *row matrix*, where we use parentheses () to denote a row matrix.

Alternatively, if $n = 1$,

$$A = \begin{bmatrix} a_{11} \\ a_{21} \\ \vdots \\ a_{m1} \end{bmatrix} = \{a_{11}, a_{21}, \dots, a_{m1}\} \quad (1-28.3)$$

contains one column. Hence it is called a *column matrix*, where, for economy of space, we use braces { } to denote a column matrix. Because the numbers $a_{11}, a_{12}, \dots, a_{1n}$ (or the numbers $a_{11}, a_{21}, \dots, a_{m1}$) may be taken as the components of a vector in n -dimensional space, it follows that a row matrix and a column matrix are sometimes called vectors of the first kind and of the second kind, respectively.

If all $a_{ij} = 0$, $1 \leq i \leq m$, $1 \leq j \leq n$, then the matrix $A = [a^{ij}] = [0]$ is called the *null matrix*.

The algebraic operations of addition, subtraction, multiplication, division, and so on of matrices are defined in terms of equivalent operations on the elements of the matrices. These operations are discussed next.

Matrix Addition. Let $A = [a_{ij}]$, $B = [b_{ij}]$, $1 \leq i \leq m$, $1 \leq j \leq n$. Then the operation of addition, denoted by $A + B$, is defined by

$$A + B = [a_{ij} + b_{ij}] = \begin{bmatrix} a_{11} + b_{11} & a_{12} + b_{12} & \dots & a_{1n} + b_{1n} \\ a_{21} + b_{21} & a_{22} + b_{22} & \dots & a_{2n} + b_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} + b_{m1} & a_{m2} + b_{m2} & \dots & a_{mn} + b_{mn} \end{bmatrix} \quad (1-28.4)$$

Matrix Subtraction. The subtraction of matrices A , B , denoted by $A - B$, as in addition, requires that A , B be of the same order. By definition,

$$A - B = [a_{ij} - b_{ij}] \quad (1-28.5)$$

If $A = B$, that is, if $a_{ij} = b_{ij}$, $A - B = [0]$, the null matrix. In other words, two matrices A , B are said to be equal if they are of the same order and their difference is the null matrix.

Multiplication of a Matrix by a Scalar. Multiplication of a matrix A by a scalar s multiplies every element a_{ij} of A by s . Thus, $sA = s[a_{ij}] = [sa_{ij}]$. Analogously, division of a matrix A by a scalar s is defined by $(1/s)A = (1/s)[a_{ij}] = [(1/s)a_{ij}]$.

Multiplication of a Matrix by a Matrix. The operation of a matrix multiplication occurs in a number of situations. For example, in Section 1-24 we found that a rotation from one set of Cartesian axes x_α to another set X_α led to the result [Eq. (1-24.11)]

$$X_\alpha = a_{\alpha\beta}x_\beta \quad \alpha, \beta = 1, 2, 3 \quad (1-28.6)$$

Similarly, a rotation from axes X_α to axes Y_α yields

$$Y_\alpha = b_{\alpha\beta}X_\beta \quad (1-28.7)$$

where $b_{\alpha\beta}$ are direction cosines between axes Y_α and X_β . Hence, substitution of Eqs. (1-28.6) into Eq. (1-28.7) yields a transformation from axes x_α directly to axes Y_α . Thus,

$$Y_\alpha = b_{\alpha\beta}a_{\beta\gamma}x_\gamma = c_{\alpha\gamma}x_\gamma \quad (1-28.8)$$

where

$$c_{\alpha\gamma} = b_{\alpha\beta}a_{\beta\gamma} \quad (1-28.9)$$

In matrix notation, we may write

$$\begin{aligned} X &= Ax \\ Y &= BX = BAx = Cx \end{aligned} \quad (1-28.10)$$

where

$$A = [a_{\beta\gamma}] \quad B = [b_{\alpha\beta}] \quad C = [c_{\alpha\gamma}] \quad (1-28.11)$$

and where summation convention holds (Section 1-23).

Generalization of Eq. (1-28.10) leads to the following definition: Given $A = [a_{ij}]$, $1 \leq i \leq m$, $1 \leq j \leq n$; $B = [b_{ij}]$, $1 \leq i \leq p$, $1 \leq j \leq q$. The product AB is defined if and only if $p = n$. When $p = n$, matrices A and B are said to be *conformable* or to *conform*. The product of two conformable matrices A (of order m by n) and B (of order n by q) is a matrix C (of order m by q), with elements c_{ij} given by the rule

$$c_{ij} = b_{i\alpha}a_{\alpha j} = b_{i1}a_{1j} + b_{i2}a_{2j} + \cdots + b_{in}a_{nj} \quad (1-28.12)$$

This rule is summarized by the following statement: The matrix C , with elements c_{ij} , is obtained by multiplication of the elements of the i th row of matrix B into the elements of the j th column of matrix A .

In general, we note that the premultiplication BA of A by B is not equal to the postmultiplication AB of A by B . Thus, in general, $BA \neq AB$. In particular, BA and AB are both defined if and only if $m = q$ and $p = n$.

More generally, the product P of p matrices (extended product) $A_1, A_2, A_3, \dots, A_p$ is defined by $P = A_1 A_2 A_3 \cdots A_p$, provided that in the order $A_1, A_2, A_3, \dots, A_p$ two adjacent matrices conform. If $A_1 = A_2 = A_3 = \cdots = A_p = A$, we obtain $P = A^p$, the p th product of A .

Square Matrices. A matrix $A = [a_{ij}]$, $1 \leq i \leq m$, $1 \leq j \leq n$ is said to be square if and only if $m = n$. A square matrix is said to be symmetric if and only if $a_{ij} = a_{ji}$. If $a_{ij} = 0$, for $i \neq j$, the matrix $A = [a_{ij}]$ is said to be a *diagonal matrix* and is denoted by $A = \text{diag}(a_{11}, a_{22}, \dots, a_{nn})$. If A is a diagonal matrix and $a_{ii} = s$ for all i , A is called a *scalar matrix*. If, in addition, $s = 1$, the matrix A consists of diagonal elements all equal to 1. Then A is called the *unit matrix* and is denoted by the symbol I ; that is, $A = I$. For any matrix B , we have $IB = BI = B$. Hence, I commutes with any matrix. Thus, the unit matrix operates on matrices in the same manner that the number 1 operates on real numbers.

Transpose of a Matrix. In operations with arrays $[a_{ij}]$, we must consider arrays $[a_{ji}]$. The matrix $[a_{ji}] = A^T$ is called the *transpose* of the matrix $A = [a_{ij}]$, and the operation of forming the transpose A^T from matrix A is called *transposition*. In particular, the transpose P^T of a product $P = AB$ of matrices A, B is $P^T = B^T A^T$, and, in like manner, the transpose of the extended product $P = A_1 A_2 \cdots A_q$ is $P^T = A_q^T A_{q-1}^T A_{q-2}^T \cdots A_1^T$.

Division of a Matrix by a Matrix. The operation of division is restricted to square matrices. Several preliminary notions are required: the concepts of the determinant $|a|$ of a square matrix $[a_{ij}]$, the cofactor A_{ij} of the elements a_{ij} in the determinant $|a|$, the adjoint matrix $\bar{A} = [A_{ji}]$ of a matrix, and the inverse of a matrix, denoted by A^{-1} .

We assume that the concept of determinant is familiar from elementary algebra. Then, for a square n by n matrix $A = [a_{ij}]$, $1 \leq i \leq n$, $1 \leq j \leq n$, we have the associated determinant $|a|$ of the matrix $[a_{ij}]$, where the number $|a|$ is defined (Birkhoff and MacLane, 2008; Lancaster and Tismenetsky, 1985; Gilbert, 2008) by

$$|a| = \begin{vmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \cdots & a_{nn} \end{vmatrix} \quad (1-28.13)$$

If $|a| \neq 0$, $A = [a_{ij}]$ is said to be nonsingular and possesses a *reciprocal* or *inverse matrix* A^{-1} such that

$$[a_{ij}]A^{-1} = AA^{-1} = I \quad (1-28.14)$$

where I is the unit matrix of the same order as A . Accordingly, the operation of matrix multiplication of a matrix by its inverse matrix is analogous to dividing a real number by itself. More generally, if B is any matrix conformable with the nonsingular inverse matrix A^{-1} ,

$$BA^{-1} = C \quad (1-28.15)$$

where C is a matrix. Equation (1-28.15) is sometimes referred to as the *division* of matrix B by matrix A . Accordingly, to divide matrix B by a conformable matrix A , we must first compute the inverse matrix A^{-1} . To compute A^{-1} , we first introduce the *adjoint matrix* $\bar{A}$ of A . The adjoint matrix $\bar{A}$ is defined by

$$\bar{A} = [A_{ji}] \quad (1-28.16)$$

where the element A_{ij} denotes the cofactor of the element a_{ij} in the determinant $|a|$ of the matrix $A = [a_{ij}]$, and $[A_{ij}]$ is the transpose of the matrix $[A_{ij}]$. The adjoint matrix $\bar{A}$ exists whether or not A is singular.

By definition of matrix multiplication and the theory of determinants, we have [Eq. (1-28.12)]

$$[a_{ij}][A_{ji}] = |a|I = S \quad (1-28.17)$$

where S is a diagonal (scalar) matrix, with $s_{ii} = |a|$, and $s_{ij} = 0$ for $i \neq j$. Dividing Eq. (1-28.17) by the determinant $|a|$ (assumed nonsingular), we obtain

$$I = \frac{[a_{ij}][A_{ji}]}{|a|} = \frac{A\bar{A}}{|a|} \quad (1-28.18)$$

Accordingly, comparison of Eqs. (1-28.14) and (1-28.18) yields

$$A^{-1} = \frac{[A_{ji}]}{|a|} = \frac{\overline{A}}{|a|} \quad (1-28.19)$$

The matrix A^{-1} is called the *inverse* or *reciprocal matrix* because of the property $AA^{-1} = I$. It plays the same role in matrix algebra as does division in ordinary algebra. Thus, if $AB = CD$ where A, B, C, D are appropriate matrices, premultiplication by A^{-1} , the inverse of A , yields $A^{-1}AB = IB = B = A^{-1}CD$.

1-29 Some Topics in the Calculus of Variations

Maxima, Minima, and Lagrange Multipliers. The problem of seeking maxima and minima of functions of several variables plays an important role in engineering. A generalization of the elementary theory of maxima and minima (or extrema) leads to the calculus of variations. For example, in the theory of extrema, consider the problem of determining for a given continuous function $f(x_1, x_2, \dots, x_n)$ of the n variables $(x_1, x_2, \dots, x_n)$ in a given region R , a point $(x_{p1}, x_{p2}, x_{p3}, \dots, x_{pn})$ at which the function f attains *maximum* or *minimum values* (i.e., *extreme values* or simply *extrema*) with respect to all points of R in a neighborhood (vicinity) of the point $(x_{p1}, x_{p2}, \dots, x_{pn})$. This problem always has a solution (point for which f is an extremum), because according to a theorem of Weierstrass, every function $f(x_1, x_2, \dots, x_n)$ that is continuous in a closed bounded region R of the variables $(x_1, x_2, \dots, x_n)$ possesses a maximum value and a minimum value in the interior of R or on the boundary of R . Analogous to the theory of a single variable, if the function $f(x_1, x_2, \dots, x_n)$ is differentiable in R and if an extreme value is attained at an interior point $P : (x_{p1}, x_{p2}, \dots, x_{pn})$, then the derivatives of f with respect to each of the x 's vanish at P . The vanishing of the derivatives of f is a necessary condition for extrema. It is not sufficient, however, as an examination of the function $f(x) = x^3$ at the point $x = 0$ shows. More generally, we define a point at which all first-order derivatives of f vanish, hence at which $df = 0$, as a stationary point S . In turn, a *stationary point* S that furnishes a maximum value or a minimum value (an extreme value) in an allowable neighborhood of S is called an *extremum*.

In some problems the choice of points $(x_1, x_2, \dots, x_n)$ is restricted to subregions of R by certain *equations of constraint* (or simply *constraints*). For example, consider the stationary values of the function $f(x_1, x_2, \dots, x_n)$ of the n variables $(x_1, x_2, \dots, x_n)$, continuous with continuous first partial derivatives, subject to the restrictions that the x 's must satisfy m equations of constraint ($m < n$)

$$g_i(x_1, x_2, \dots, x_n) = 0 \quad i = 1, 2, \dots, m \quad (1-29.1)$$

The direct approach to this problem is to eliminate m of the variables from f by means of Eq. (1-29.1). Then seek the stationary values of $f(y_1, y_2, \dots, y_{n-m})$, where $y_1, y_2, \dots, y_{n-m}$ denote the remaining $n - m$ variables. However, because

the elimination of any m of the variables is arbitrary and the process of elimination from Eq. (1-29.1) may be nontrivial, an alternative approach attributed to Lagrange is often employed (the *Lagrange multiplier method*). This method has the advantage of retaining symmetry in the calculations (arbitrary elimination of m variables is avoided) and of routine elegance.

Lagrange's method of multipliers consists of forming a new function F such that

$$F(x_1, x_2, \dots, x_n; \lambda_0, \lambda_1, \dots, \lambda_m) = \lambda_0 f(x_1, x_2, \dots, x_n) + \sum_{i=1}^m \lambda_i g_i(x_1, x_2, \dots, x_n) \quad (1-29.2)$$

where the $\lambda_i, i = 0, 1, 2, \dots, m$, are called the *Lagrange multipliers*. Then stationary values of F are sought over the unrestricted range of the variables $(x_1, x_2, \dots, x_n)$ by the requirements

$$\begin{aligned} \frac{\partial F}{\partial x_1} = 0, \quad \frac{\partial F}{\partial x_2} = 0, \dots, \quad \frac{\partial F}{\partial x_n} = 0 \\ \frac{\partial F}{\partial \lambda_1} = g_1 = 0, \quad \frac{\partial F}{\partial \lambda_2} = g_2 = 0, \dots, \quad \frac{\partial F}{\partial \lambda_m} = g_m = 0 \end{aligned} \quad (1-29.3)$$

These equations suffice to determine the stationary points $(x_{p1}, x_{p2}, \dots, x_{pn})$ and the Lagrange multipliers $\lambda_1, \lambda_2, \dots, \lambda_m$. Because F is homogeneous in the λ 's [Eq. (1-29.2)], we may take $\lambda_0 = 1$.

Equations (1-29.3) show that the stationary points for F are the same as the stationary values of f subject to the constraints of Eq. (1-29.1). The Lagrange multiplier method is useful in the theory of principal values of stress and strain (Chapters 2 and 3).

More generally, the above results may be summarized as follows⁷:

Given a function $f(x_1, x_2, \dots, x_n)$ of n variables $(x_1, x_2, \dots, x_n)$ subject to m constraints $g_i(x_1, x_2, \dots, x_n) = 0, i = 1, 2, \dots, m$. Let f and g_i possess continuous first partial derivatives in a region R of the x space. Furthermore, let the Jacobian J be nonzero; that is,

$$J = \frac{\partial(g_1, g_2, \dots, g_m)}{\partial(a_1, a_2, \dots, a_m)} \neq 0 \quad (1-29.4)$$

where the set of variables $(a_1, a_2, \dots, a_m)$ is some selection of m variables from the extremum $(x_{p1}, x_{p2}, \dots, x_{pn})$. Then the stationary values of f subjected to

⁷For an analytical proof of the Lagrange multiplier method, see Courant (1992), pp. 192–199 (footnote 5).

the constraints $g_i = 0, i = 1, 2, \dots, m$ are identical to the stationary values of the function

$$F(x_1, x_2, \dots, x_n; \lambda_1, \lambda_2, \dots, \lambda_m) = f(x_1, x_2, \dots, x_n) + \sum_{i=1}^m \lambda_i g_i(x_1, x_2, \dots, x_n) \quad (1-29.5)$$

In cases where the constraints $g_i = 0$ are algebraic relations, the Lagrange multipliers are constant parameters. However, more generally (Langhaar, 1989), when the equations of constraint require the x_i to be solutions of differential equations, the Lagrange multipliers may be functions of one or more of the variables x_i .

Variation of a Function. First Variation of an Integral. Stationary Value of an Integral. As in the theory of ordinary maxima and minima, the calculus of variations is concerned with the problem of extreme values (stationary values). However, in contrast to the ordinary extremum problem of a function of a finite number of independent variables, the calculus of variations deals with functions of functions, or simply functionals (Courant and Hilbert, 1989).

The simplest type of problem in the calculus of variations may be outlined as follows: Let $F(x, y, y')$ be a given function of the three arguments x, y, y' that is continuous and has continuous first and second derivatives in the region of the arguments. Because F is a function of x , an integral

$$I(y) = \int_{x_0}^{x_1} F(x, y, y') dx$$

becomes a definite number depending upon the behavior of the function $y = y(x)$, the *argument function*. That is, the integral $I(y)$ becomes a function of the argument function $y(x)$ or, in other words, a *functional*. The fundamental problem of the calculus of variations may be stated in this form: Among all functions $y = y(x)$ that are defined and continuous and possess continuous first and second derivatives in the interval $x_0 \leq x \leq x_1$ and for which boundary values $y_0 = y(x_0), y_1 = y(x_1)$ are given, determine that function $y = u(x)$ for which the integral $I(y)$ has the smallest possible value (or the largest possible value). The conditions imposed upon the argument function $y(x)$ are called *conditions of admissibility*, and we speak of argument functions that satisfy the conditions of admissibility as *admissible functions*. The admissible functions $y(x)$ form a class C. In the above formulations we required that $y(x)$ be continuous with its first and second derivatives. Actually, the existence of $I(y)$ requires only that F , hence $y'(x)$, be sectionally continuous. The more restricted admissible conditions limit the class C in which functions $y(x)$ are sought. However, it may be shown that the function $y = f(x)$, which minimizes I when the broader class of admissible conditions is allowed, always lies in the more restricted class of admissible functions (Courant and Hilbert, 1989).

Accordingly, our objective is to determine *necessary conditions* that an admissible function $y = u(x)$ gives a maximum or minimum value (extreme value) to the

integral $I(y)$. The method employed is analogous to that of the extreme problem of determining the extreme value of a function of a single variable. Thus, we assume that $y = u(x)$ is the solution, say, a minimum. [The problem of determining a maximum may be dispensed with, as the method of seeking a maximum is the same as that for seeking a minimum with F replaced by $-F$ in $I(y)$.] Then for any other admissible function the value of I must increase. Because we seek necessary conditions, it suffices to consider admissible functions that lie infinitesimally close to the solution $y = u(x)$. Hence, we consider the class of admissible functions

$$\bar{y} = y(x) + \varepsilon \eta(x) = y(x) + \delta y$$

where ε is a parameter and $\eta(x)$ is a function in the class of admissible functions (i.e., has continuous first and second derivatives in $x_0 \leq x \leq x_1$ and vanishes at $x = x_0$ and $x = x_1$). The quantity $\delta y = \varepsilon \eta(x)$ is called the variation of the function $y(x)$. Then if ε is sufficiently small, the admissible functions $\bar{y}$ lie in an arbitrarily small neighborhood of the extremum $y = u(x)$. Hence, the integral $J = I(y + \varepsilon \eta)$ may be regarded as a function of ε , which must attain a minimum at $\varepsilon = 0$ relative to all values of ε in a sufficiently small neighborhood of $\varepsilon = 0$. Consequently, $dJ/d\varepsilon|_{\varepsilon=0} = J'(0) = 0$ is a necessary condition that $I(y)$ attain a minimum for $y = u(x)$. More generally, without regard to maximum or minimum, we say that the integral I is *stationary* for $y = u(x)$. Thus, with $J(\varepsilon) = I(y + \varepsilon \eta) = \int_{x_0}^{x_1} F(x, y + \varepsilon \eta, y' + \varepsilon \eta')/dx$, differentiation yields the necessary condition

$$J'(0) = \int_{x_0}^{x_1} (F_y \eta + F_{y'} \eta') dx = 0$$

that $I(y)$ be stationary for all admissible $\eta(x)$. Integration by parts and use of the conditions $\eta(x_0) = \eta(x_1) = 0$ yield

$$J'(0) = \int_{x_0}^{x_1} \eta \left(F_y - \frac{d}{dx} F_{y'} \right) dx = 0$$

which must hold for arbitrary admissible functions η . Hence, by the fundamental theorem of the calculus of variations,⁸

$$F_y - \frac{dF_{y'}}{dx} = 0 \quad (1-29.6)$$

Equation (1-29.6) is the *Euler differential equation* for the integral $I(y)$. It is a necessary condition that $I(y)$ possess a *stationary value*.

Recalling the definition $\delta y = \varepsilon \eta(x)$, and noting that $\eta(x) = d\bar{y}/d\varepsilon$, we may interpret the symbol δ to denote the differential obtained when ε is regarded as the

⁸See Langhaar (1989). Langhaar gives an elegant approach to the derivation of the Euler equation in Section 3-2.

independent variable. Then the equation

$$\delta I = \eta \left. \frac{dJ}{d\varepsilon} \right|_{\varepsilon=0} = \int_{x_0}^{x_1} (\eta F_y + \eta' F_{y'}) dx \quad (1-29.7)$$

is called the *first variation* of the integral I . Hence, the terminology *stationary character of an integral* means the same thing as *vanishing of the first variation* of the integral.

REFERENCES

- Abraham, F., Walkup, R., Gao, H., Duchaineau, M., La Rubia, T. D., and Seager, M. 2002. Simulating Materials Failure by Using up to One Billion Atoms and the World's Fastest Computer: Work-hardening, *Proc. Natl. Acad. Sci. USA*, 99(9): 5783–5787.
- Arroyo, M., and Belytschko, T. 2005. Continuum Mechanics Modeling and Simulation of Carbon Nanotubes, *Mechanica*, 40: 455–469.
- Atrek, E., Gallagher, R. H., Ragsdell, K. M., and Zienkiewicz, O. C. (eds.) 1984. *New Directions in Optimal Structural Design*. New York: John Wiley & Sons.
- Belytschko, T., and Xiao, S. P. 2003. Coupling Methods for Continuum Model with Molecular Model, *Intl. J. Multiscale Comput. Engr.*, 1(1): 115–126.
- Belytschko, T., Xiao, S. P., Schatz, G. C., and Ruoff, R. S. 2002. Atomistic Simulations of Nanotubes Fracture, *Phys. Rev. B*, 65: 235–430.
- Birkhoff, G., and MacLane, S. 2008. *A Survey of Modern Algebra*, Natick, MA: A.K. Peters Ltd.
- Boresi, A. P., Chong, K. P., and Saigal, S. 2002. *Approximate Solution Methods in Engineering Mechanics*. New York: John Wiley & Sons.
- Brebbia, C. A. 1988. *The Boundary Element Method for Engineers*. New York: John Wiley & Sons.
- Cheung, Y. K., and Tham, L. G. 1997. *Finite Strip Method*. Boca Raton, FL: CRC Press.
- Chong, K. P. 2004. Nanoscience and Engineering in Mechanics and Materials, *J. Phys. Chem. Solids*, 65: 1501–1506.
- Chong, K. P. 2010. Translational Research in Nano, Bio Science and Engineering, 1st Global Congress on Nanoengineering for Medicine & Biology, ASME Proc. NEMB 2010–13384, Houston, TX.
- Chong, K. P., and Davis, D. C. 2000. Engineering System Research in the Information Technology Age, *ASCE J. Infrastruct. Syst.*, March: 1–3.
- Chong, K. P., and Smith, J. W. 1984. Mechanical Characterization, in Chong, K. P., and Smith, J. W. (eds.), *Mechanics of Oil Shale*, Chapter 5. London: Elsevier Applied Science Publishing Company.
- Chong, K. P., Liu, S. C., and Li, J. C. (eds.) 1990. *Intelligent Structures*. London: Elsevier Applied Science Publishing Company.
- Chong, K. P., Dillon, O. W., Scalzi, J. B., and Spitzig, W. A. 1994. Engineering Research in Composite and Smart Structures, *Compos. Engr.*, 4(8): 829–852.
- Courant, R. 1992. *Differential and Integral Calculus*, Vols. I and II. New York: John Wiley & Sons.

- Courant, R., and Hilbert, D. 1989. *Methods of Mathematical Physics*, Vols. I and II. New York: John Wiley & Sons.
- Dally, J. W., and Riley, W. F. 2005. *Experimental Stress Analysis*, Knoxville, TN: College House Enterprises.
- Dove, R. C., and Adams, P. H. 1964. *Experimental Stress Analysis and Motion Measurement*. Columbus, OH: Charles E. Merrill Publishing Company.
- Dvorak, G. J. (ed.) 1999. Research Trends in Solid Mechanics, *Intl. J. Solids Struct.*, 37(1&2): Special Issues.
- Eisenhart, L. P. 2005. *Coordinate Geometry*. New York: Dover Publishing.
- Ellis, T. M. R., and Semenov, O. I. (eds.) 1983. *Advances in CAD/CAM*. Amsterdam: North-Holland Publishing Company.
- Eringen, A. C. 1980. *Mechanics of Continua*. Melbourne, FL: Krieger Publishing Company.
- Fosdick, L. D. (ed.) 1996. *An Introduction to High-Performance Scientific Computing*. Boston: MIT Press.
- Fung, Y. C. 1967. Elasticity of Soft Tissues in Simple Elongation, *Am. J. Physiol.*, 28: 1532–1544.
- Fung, Y. C. 1983. On the Foundations of Biomechanics, *ASME J. Appl. Mech.*, 50: 1003–1009.
- Fung, Y. C. 1990. *Biomechanics: Motion, Flow, Stress, and Growth*. New York: Springer.
- Fung, Y. C. 1993. *Biomechanics: Material Properties of Living Tissues*. New York: Springer.
- Fung, Y. C. 1995. Stress, Strain, Growth, and Remodeling of Living Organisms, *Z. Angew. Math. Phys.*, 46: S469–482.
- Gilbert, L. 2008. *Elements of Modern Algebra*. Belmont, CA: Brooks/Cole Publishing.
- Goursat, E. 2005. *A Course in Mathematical Analysis*, Vol. I, Section 149. Ann Arbor, MI: Scholarly Publishing Office, University of Michigan Library.
- Green, A. E., and Adkins, J. E. 1970. *Large Elastic Deformations*, 2nd ed. London: Oxford University Press.
- Green, A. E., and Zerna, W. 2002. *Theoretical Elasticity*. New York: Dover Publications.
- Greenspan, D. 1965. *Introductory Numerical Analysis of Elliptic Boundary Value Problems*. New York: Harper & Row.
- Hildebrand, F. B. 1992. *Methods of Applied Mathematics*. New York: Dover Publications.
- Humphrey, J. D. 2002. Continuum Biomechanics of Soft Biological Tissues, *Proc. R. Soc. Lond. A*, 459: 3–46.
- Ince, E. L. 2009. *Ordinary Differential Equations*. New York: Dover Publications.
- Khang, D.-Y., Jiang, H., Huang, Y., and Rogers, J. A. 2006. A Stretchable Form of Single-Crystal Silicon for High-Performance Electronics on Rubber Substrates, *Science*, 311: 208–212.
- Kirsch, U. 1993. *Structural Optimization*. New York: Springer.
- Knops, R. J., and Payne, L. E. 1971. *Uniqueness Theorems in Linear Elasticity*. New York: Springer.
- Lamit, L. 2007. *Moving from 2D to 3D for Engineering Design: Challenges and Opportunities*. BookSurge, LLC; www.booksurge.com.

- Lancaster, P., and Tismenetsky, M. 1985. *The Theory of Matrices*, 2nd ed. New York: Academic Press.
- Langhaar, H. L. 1989. *Energy Methods in Applied Mechanics*. Melbourne, FL: Krieger Publishing Company.
- Liu, W. K., Karpov, E. G., Zhang, S., and Park, H. S. 2004. An Introduction to Computational Nanomechanics and Materials, *Comput. Methods Appl. Mech. Engr.*, 193: 1529–1578.
- Londer, R. 1985. Access to Supercomputers. *Mosaic* 16(3): 26–32.
- Love, A. E. H. 2009. *A Treatise on the Mathematical Theory of Elasticity*. Bel Air, CA: BiblioBazaar Publishing.
- Masud, A., and Kannan, R. 2009. A Multiscale Framework for Computational Nanomechanics: Application to the Modeling of Carbon Nanotubes, *Intl. J. Numer. Meth. Engr.*, 78(7): 863–882.
- Mendelson, A. 1983. *Plasticity: Theory and Applications*. Melbourne, FL: Krieger Publishing Company.
- Meyers, M. A., Chen, P., Lin, A. Y., and Seki, Y. 2008. Biological Materials: Structure and Mechanical Properties, *Prog. Mat. Sci.*, 53: 1–206.
- Moon, F. C., et al. 2003. *Future Research Directions in Solid Mechanics*. Final Report of the American Academy of Mechanics (AAM), Submitted to the National Science Foundation (Program Director: K. P. Chong); *AAM Mechanics*, Vol. 32, Nos. 7–8.
- Morris, M., and Brown, O. E. 1964. *Differential Equations*, 4th ed. Englewood Cliffs, NJ: Prentice-Hall.
- Muskhelishvili, N. I. 1975. *Some Basic Problems of the Mathematical Theory of Elasticity*. Leyden, The Netherlands: Noordhoff International Publishing Company.
- Naghdi, P. M., and Hsu, C. S. 1961. On a Representation of Displacements in Linear Elasticity in Terms of Three Stress Functions. *J. Math. Mech.*, 10(2): 233–246.
- Oden, J. T. (ed.) 2000. *Research Directions in Computational Mechanics*, sponsored by the United States Committee in Theoretical and Applied Mechanics, NRC Publ.
- Oden, J. T. (Chair), 2006. Simulation-Based Engineering Science: Revolutionizing Engineering Science Through Simulation. *Report of the National Science Foundation Blue Ribbon Panel on Simulation-Based Engineering Science*, available at http://www.nsf.gov/publications/pub_summ.jsp?ods_key=sbes0506.
- Pipes, L. 1959. *Applied Mathematics for Engineers and Physicists*, 2nd ed. New York: McGraw-Hill Book Company.
- Reed, M. A., and Kirk, W. P. (eds.). 1989. *Nanostructure Physics and Fabrication*. New York: Academic Press.
- Ritchie, R. O., Buehler, M. J., and Hansma, P. 2009. Plasticity and Toughness in Bone, *Phys. Today*, June: 41–47.
- Rogers, C. A., and Rogers, R. C. (eds.) 1992. *Recent Advances in Adaptive and Sensory Materials*. Lancaster, PA: Technomic Publishing.
- Ruud, C. O., and Green, R. E., Jr. 1984. *Nondestructive Methods for Material Property Determination*. New York: Plenum Press.
- Schreiber, E., Orson, L. A., and Soga, N. 1973. *Elastic Constants and Their Measurements*. New York: McGraw-Hill Book Company.
- Spain, B. 2003. *Tensor Calculus*. New York: Dover Publishing.

- Srivastava, D., Makeev, M. A., Menon, M., and Osman, M. 2007. Computational Nanomechanics and Thermal Transport in Nanotubes and Nanowires, *J. Nanosci. Nanotech.*, 8(1): 1–23.
- Sternberg, E. 1960. On Some Recent Developments in the Linear Theory of Elasticity, in *Structural Mechanics*. Elmsford, NY: Pergamon Press.
- Stippes, M. 1967. A Note on Stress Functions, *Intl. J. Solids Struct.*, 3: 705–711.
- Synge, J. L., and Schild, A. 1978. *Tensor Calculus*. New York: Dover Publications.
- Thoft-Christensen, P., and Baker, M. J. 1982. *Structural Reliability Theory and Its Applications*. Berlin: Springer.
- Timp, G. (ed.). 1999. *Nanotechnology*. New York: Springer.
- Trimmer, W. (ed.) 1990. *Micromechanics and MEMS*. Piscataway, NJ: IEEE Press.
- Tsomanakis, Y., Lagaros, N. D., and Papadrakakis, M. (eds.) 2008. *Structural Design Optimization Considering Uncertainties*. London: Taylor & Francis.
- Udd, E. (ed.). 1995. *Fibre Optic Smart Structures*, New York: John Wiley & Sons.
- Wagner, G. J., Jones, R. E., Templeton, J. A., and Parks, M. L. 2008. An Atomistic-to-Continuum Coupling Method for Heat Transfer in Solids, *Comput. Methods Appl. Mech. Engr.*, 197: 3351–3365.
- Wen, Y.-K. (ed.) 1984. *Probabilistic Mechanics and Structural Reliability*. New York: American Society of Civil Engineers.
- Yang, H. Y., and Chong, K. P. 1984. Finite Strip Method with X-Spline Functions. *Comput. Struct.*, 18(1): 127–132.
- Yang, L. T., and Pan, Y., (eds.) 2004. *High Performance Scientific and Engineering Computing: Hardware and Software Support*. Norwell, MA: Kluwer Academic Publisher.
- Yao, J. T. P. 1985. *Safety and Reliability of Existing Structures*. Boston: Pitman Advanced Publishing Program.
- Zienkiewicz, O. C., and Taylor, R. L. 2005. *The Finite Element Method*, 6th ed., London: Elsevier Butterworth-Heinemann Publishing.

BIBLIOGRAPHY

- Boresi, A. P., and Schmidt, R. J. 2000. *Engineering Mechanics: Dynamics*, Pacific Groove, CA: Brooks/Cole Publishing.
- Brand, L. *Vector and Tensor Analysis*. New York: John Wiley & Sons, 1962.
- Chong, K. P., Dewey, B. R., and Pell, K. M. *University Programs in Computer-Aided Engineering, Design, and Manufacturing*. Reston, VA: ASCE, 1989.
- Danielson, D. A. *Vectors and Tensors in Engineering and Physics*. Boulder, CO: Westview Press, 2003.
- Dym, C. L., and Shames, I. H. *Solid Mechanics: A Variation Approach*. New York: McGraw-Hill Book Company, 1973.
- Edelen, D. G. B., and Kydonieffs, A. D. *An Introduction to Linear Algebra for Science and Engineering*, 2nd ed. New York: Elsevier Science Publishing Company, 1980.
- Eisele, J. A., and Mason, R. M. *Applied Matrix and Tensor Analysis*. New York: Wiley-Interscience Publishers, 1970.
- Eisberg, R., and Resnick, R. *Quantum Physics*, New York: John Wiley & Sons, 1985.

- Gere, J. M., and Weaver, W. *Matrix Algebra for Engineers*, 2nd ed. Boston: Pringle, Weber, and Scott Publishers, 1984.
- Jeffreys, H. *Cartesian Tensors*. New York: Cambridge University Press, 1987.
- Kemmer, N. *Vector Analysis*. London: Cambridge University Press, 1977.
- Lur , A. I. *Three-Dimensional Problems of the Theory of Elasticity*. New York: Wiley-Interscience Publishers, 1964.
- Nash, W. A. *The Mathematics of Nonlinear Mechanics*. Boca Raton, FL: CRC Press, 1993.
- Nemat-Nasser, S., and Hori, M. *Micromechanics*, 2nd ed. Amsterdam: North-Holland, 1999.
- Poole, Jr., C. P., and Owens, F. J. *Introduction to Nanotechnology*. Hoboken, NJ: Wiley-Interscience, 2003.
- Roco, M. C. The New Engineering World: A National Investment Strategy Is Key to Transforming Nanotechnology from Science Fiction to an Everyday Engineering Tool. *Mechanical Engineering-CIME*. American Society of Mechanical Engineers, 127 (4), S6(6), 2005.
- Roco, M. C. National Nanotechnology Initiative-Past, Present, Future. *Handbook on Nanoscience, Engineering and Technology*, 2nd ed., London: Taylor and Francis, pp. 3.1-3.26, 2007.
- Ting, T. C. T. *Anisotropic Elasticity*. New York: Oxford University Press, 1996.