

1 Conduction Heat Transfer and Formulation

1.1 INTRODUCTION

There are many practical engineering problems that require the analysis of problems involving the transfer of heat. The solution of the equation of heat conduction is sufficient in many cases, while the area of application expands considerably if the equations governing coupled heat and mass transfer and/or thermal convection are considered. Combining a thermal analysis procedure with a thermal stress predictive capability can provide answers to questions of immediate concern in many industrial processes.

In the field of civil engineering, there are numerous examples of practical problems where the behaviour of the system under consideration may be predicted via a heat transfer analysis. The thermal properties of materials used for construction purposes often necessitate the use of some form of insulation to counter the effects of the wide range of temperatures produced by seasonal climatic changes. This is particularly true in highway construction, where it is important to predict the placement and quantity of insulating material in order to prevent possible damage to the road surface due to frost heave. At excavation sites, where the soil may be unstable because of friability or water saturation, artificial freezing of the soil is often employed to render it structurally stable. This typically involves the circulation of a refrigerated liquid through pipes sunk in concentric circular patterns around the excavation site. The high cost of the process, in terms of the refrigeration plant and energy consumption, means that a numerical simulation of the process can prove to be extremely useful. If the numerical method makes accurate predictions of the advance of the frozen region, the numerical results can be used to aid the design of the optimal arrangement of the cooling pipes and to predict the required running time necessary to produce an ice wall of the required thickness. In the foundry, it is imperative to understand the process of solidification in metal casting and the effect this produces on porosity formation. In high pressure die casting, it is often required to design a cooling system that will ensure the desired surface finish. Problems of this type can be investigated by a heat conduction analysis, taking account

of latent heat effects and allowing for the variation in the thermal properties of the materials with temperature.

For the simulation of problems involving porous capillary materials, such as timber or ceramics, a coupled heat and mass transfer analysis is required. The optimal drying rate to ensure a reasonably stress-free end product will be tied in with the energy cost of the kiln schedule. Mathematical models, which can accurately predict the movement of moisture and heat in porous materials, are often the only means of gaining a better insight into the physical process. Similar models prove useful in designing geothermal energy extraction systems or in the analysis of advanced oil recovery by thermal methods. By increasing the complexity of the governing equations, it is possible to predict the aerodynamic heating of structures, such as re-entry vehicles, or of turbine blades in a jet engine. A related problem here is the computation of the thermally induced stresses and the prediction of the working life of the component. Such information is essential in the safety analysis of nuclear reactors and also in the modern steel industry, where the production rate of continuously cast steel is limited by the cracking induced by thermal stress development in the strand.

Exact analytical solutions of the governing equations of heat transfer can only be obtained for problems in which restrictive simplifying assumptions have been made with respect to geometry, material properties and boundary conditions. There is therefore no option but to turn to numerical solution methods for the analysis of practical problems, where such simplifications are not generally possible. The finite element method, with its flexibility in dealing with complex geometries, is an ideal approach to employ in the solution of such problems. However, the newcomer frequently finds the gap between finite element theory and practice quite daunting and it is the objective of this book to attempt to bridge this gap. The material included proceeds in an orderly fashion from the governing differential equations to the finite element formulation.

Initially, the basic differential equation governing heat conduction will be derived and the concept of a variational formulation of the problem will be introduced. The finite element method will then be used to generate a system of simultaneous equations which have to be solved to obtain the approximation to the temperature field in the body of interest. The process will be demonstrated in detail for the simple case of linear, steady-state problems in both one and two dimensions. For the solution of time-dependent (transient) problems, finite difference methods will be employed to determine the variation of the solution with time. Different time marching algorithms are presented and their relative merits discussed.

When the variation of the thermal properties of the materials are included in the mathematical model, the resulting equations are non-linear and techniques for dealing with this added complication are described. Some examples for which exact solutions are available are also given for one-, two- and three-dimensional problems. Solution of the non-linear problems can be extended to include the effects of latent heat addition or removal as in melting and solidifica-

tion problems. These types of problems, along with some practical examples, are dealt with in the chapter on phase change. Convection heat transfer is important when the heat exchange is between a solid and a fluid. The solution methods for such problems and the difficulties encountered in the application of the standard Galerkin method are also dealt with. As has been noted above, there are many practical problems where the determination of the temperature (or moisture) distribution is only required as a prerequisite for the calculation of the resulting stress field. Methods for the determination of thermal or drying stresses in such problems will be detailed.

1.2 MODELLING OF HEAT CONDUCTION

1.2.1 Derivation of the governing equation

The equation governing the conduction of heat in a continuous medium can be derived by imposing the principle of conservation of heat energy over an arbitrary fixed volume, V , of the medium which is bounded by a closed surface S . For convenience the conservation statement is expressed in rate form and is written as:

$$\begin{aligned} \text{rate of increase of heat in } V = & \text{rate of heat conduction into } V \text{ across } S \\ & + \text{rate of heat generation within } V \end{aligned} \quad (1.2.1)$$

If u denotes the specific internal energy of the medium, then

$$\text{rate of increase of heat in } V = \int_V \rho \frac{\partial u}{\partial t} dV \quad (1.2.2)$$

where ρ is the density of the medium. Introducing the specific heat, c , of the medium defined by

$$c = \frac{du}{dT} \quad (1.2.3)$$

where T is the temperature, means that we can write equation (1.2.2) as

$$\text{rate of increase of heat in } V = \int_V \rho c \frac{\partial T}{\partial t} dV \quad (1.2.4)$$

To obtain an expression for the rate at which heat is conducted into V across S , we make use of Fourier's Law of Conduction. This is an empirical relationship which states that, for a surface with unit normal vector $\mathbf{n}$, the rate at which heat is conducted across the surface, per unit area, in the direction of $\mathbf{n}$ is given by

$$q = -k(\text{grad } T) \cdot \mathbf{n} = -k \frac{\partial T}{\partial n} \quad (1.2.5)$$

where k is a property of the medium termed the thermal conductivity. In this equation $\partial/\partial n$ denotes differentiation in the direction of $\mathbf{n}$ and q is termed the flux of heat in this direction. Thus, if $\mathbf{n}$ denotes the outward unit normal to S , it follows that

$$\begin{aligned} & \text{rate of heat conduction into } V \text{ across } S \\ &= \int_S -q \, dS - \int_S k(\text{grad } T) \cdot \mathbf{n} \, dS - \int_V \text{div}(k \text{ grad } T) \, dV \end{aligned} \quad (1.2.6)$$

where the Divergence Theorem has been applied. If it is assumed that heat generation in the medium is occurring at a rate Q per unit volume, then

$$\text{rate of heat generation within } V = \int_V Q \, dV \quad (1.2.7)$$

Using equations (1.2.4), (1.2.6) and (1.2.7) in (1.2.1) produces the conservation statement:

$$\int_V \left(\rho c \frac{\partial T}{\partial t} - \text{div}(k \text{ grad } T) - Q \right) dV = 0 \quad (1.2.8)$$

and, since the volume V was arbitrarily chosen initially, it follows that

$$\rho c \frac{\partial T}{\partial t} = \text{div}(k \text{ grad } T) + Q \quad (1.2.9)$$

everywhere in the medium. This is the familiar form of the heat conduction equation for a non-stationary system.

If the medium is anisotropic, i.e. the conductivity depends upon the direction, the form of the heat conduction equation is modified to

$$\rho c \frac{\partial T}{\partial t} = \text{div}(\mathbf{k} \text{ grad } T) + Q \quad (1.2.10)$$

where

$$\mathbf{k} = \begin{bmatrix} k_{xx} & k_{xy} & k_{xz} \\ k_{yx} & k_{yy} & k_{yz} \\ k_{zx} & k_{zy} & k_{zz} \end{bmatrix} \quad (1.2.11)$$

is a conductivity tensor and, for example, k_{xy} denotes the thermal conductivity in the x direction across a surface with normal in the y direction. If the conductivity k and the specific heat capacity ρc are assumed to be constant, and if the heat generation rate Q is independent of T , then equation (1.2.9) is linear and can be written as

$$\frac{1}{\alpha_1} \frac{\partial T}{\partial t} = \nabla^2 T + \frac{Q}{k} \quad (1.2.12)$$

where $\alpha_1 = k/\rho c$ is termed the thermal diffusivity of the medium and ∇^2 denotes

the Laplacian operator defined, in Cartesian coordinates, by

$$\nabla^2 = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2} \quad (1.2.13)$$

In the absence of heat generation within the medium, equation (1.2.12) reduces to the standard diffusion equation:

$$\frac{1}{\alpha_1} \frac{\partial T}{\partial t} = \nabla^2 T \quad (1.2.14)$$

If, in addition, the temperature does not vary with time, steady-state conditions are said to exist and, in this case, the governing equation simplifies further to

$$\nabla^2 T = 0 \quad (1.2.15)$$

which is just the Laplace equation.

1.2.2 Initial and boundary conditions

Suppose that the solution of the heat conduction equation (1.2.9) is required over an arbitrary domain Ω bounded by a closed surface, Γ , as illustrated in Figure 1.2.1. If the problem being modelled is independent of time (i.e. steady), the solution will be uniquely defined provided that we are able to supply appropriate boundary conditions. For the steady heat conduction equation, one condition has to be specified at each point of the boundary curve Γ and typical conditions of practical interest would be:

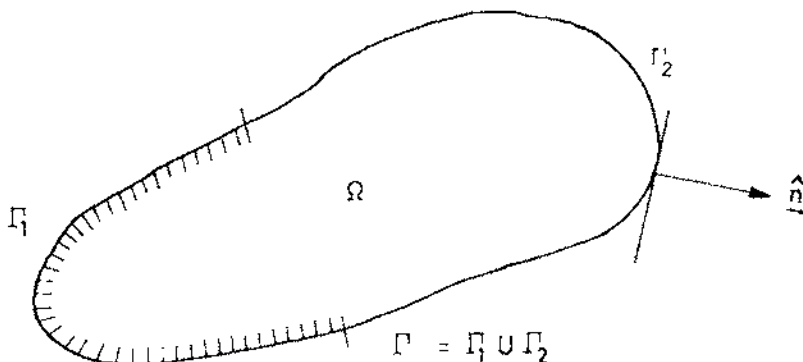


Figure 1.2.1 General domain and boundary

(a) the value of the temperature is prescribed, e.g.

$$T = f(\mathbf{x}) \quad \text{for all } \mathbf{x} \text{ on } \Gamma_1 \quad (1.2.16)$$

or

(b) the value of the outward normal heat flux is prescribed, e.g.

$$q = -k(\text{grad } T) \cdot \mathbf{n} = -k \frac{\partial T}{\partial n} = \mathcal{N}(\mathbf{x}, T) + \mathcal{N}_c(\mathbf{x}, T) + \mathcal{N}_r(\mathbf{x}, T) \quad \text{for all } \mathbf{x} \text{ on } \Gamma_2 \quad (1.2.17)$$

Here $f, \mathcal{N}, \mathcal{N}_c$ and $\mathcal{N}_r$ are prescribed functions of $\mathbf{x}$ and T and $\Gamma = \Gamma_1 \cup \Gamma_2$, $\Gamma_1 \cap \Gamma_2 = 0$. In equation (1.2.17) $\mathcal{N}$ denotes a specified heat flux; $\mathcal{N}_c$ denotes a convective heat flux, defined as

$$\mathcal{N}_c = \alpha(T - T_\infty) \quad (1.2.18)$$

where α is a coefficient of surface heat transfer and T_∞ is the specified ambient temperature of the surrounding medium; $\mathcal{N}_r$ is a radiative heat flux which is defined as

$$\mathcal{N}_r = \epsilon \sigma (T^4 - T_\infty^4) \quad (1.2.19)$$

where σ is the Stefan-Boltzman constant and ϵ is the emissivity of the surface, defined as the ratio of the heat emitted by the surface to the heat emitted by a black body at the same temperature.

When the problem being modelled is time dependent (transient), the solution is uniquely determined provided that an initial condition is given together with a boundary condition at each point of the boundary Γ of the domain. The initial condition should give the distribution of the temperature over the entire region Ω at an initial time, usually taken to be the time $t = 0$. In addition, in a transient problem, the functions $f, \mathcal{N}, \mathcal{N}_c$ and $\mathcal{N}_r$ of equations (1.2.16) and (1.2.17) may vary with time.

1.2.3 Interface conditions

When the region Ω , consists of two or more different materials, as shown in Figure 1.2.2, boundary conditions of continuity of temperature and flux across material interfaces have to be applied in order to uniquely define the solution. If we consider an interface $\tilde{\Gamma}$ between two materials, designated Ω_1 and Ω_2 , with unit normal vector $\hat{\mathbf{n}}$ in the direction into Ω_1 , these conditions can be written in the form

$$\left. \begin{aligned} T_1(\mathbf{x}, t) &= T_2(\mathbf{x}, t) \\ \hat{\mathbf{n}} \cdot (k \text{ grad } T)_1 &= \hat{\mathbf{n}} \cdot (k \text{ grad } T)_2 \end{aligned} \right\} \quad \text{for all } \mathbf{x} \text{ on } \tilde{\Gamma} \text{ and all } t > 0 \quad (1.2.20)$$

where the subscripts 1 and 2 denote the conditions appropriate to regions Ω_1 and Ω_2 , respectively.

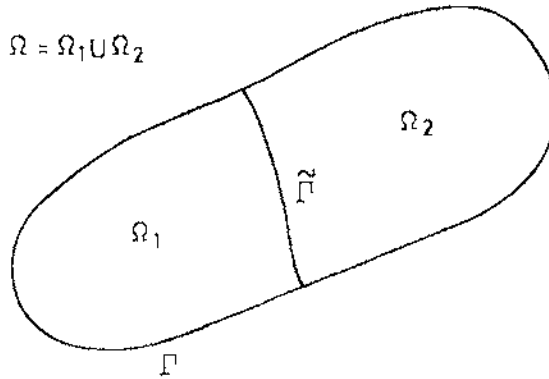


Figure 1.2.2 Composite material domain and interface.

1.2.4 Initial, boundary and interface conditions in a practical example

An example of a realistic physical problem is now used to illustrate how the correct boundary and initial conditions may be identified in practice. A tapered slab of aluminium bronze is cast in a resin bonded silica sand mould and a thermal analysis is to be made of the process. This problem involves the added complication that the metal is initially liquid and undergoes a change of phase during the transient. Methods of handling phase change problems will be introduced in Chapter 5 and the complications associated with the phase change process can be ignored for the purposes of this illustration. To reduce computational costs, the decision is made to undertake a two-dimensional analysis, and a vertical cross-section through the mould and casting, shown in Figure 1.2.3, is examined. In the actual problem, metal at 1156°C is poured into the mould through the riser entrance CD. The simulation of the pouring stage can be expected to pose additional difficult problems and so the conduction analysis will start from the time when the mould is full, with the metal taken to be at a uniform initial temperature of 1156°C . The external boundary EFGHAB of the sand will be assumed to be maintained at a constant temperature of 20°C while the initial temperature in the sand is also taken to be at this value. The upper boundary, CD, of the riser will be held constant at a temperature 1156°C , while the sides BC and DE of the riser will be assumed to be perfectly insulated and therefore subjected to a boundary condition of zero normal heat flux. The initial temperature at the interface between the mould and the metal is taken to be 1080°C , which is the liquidus temperature of the metal and perfect conduction is assumed at the interface during the transient so that the boundary conditions of equation (1.2.20) may be applied.

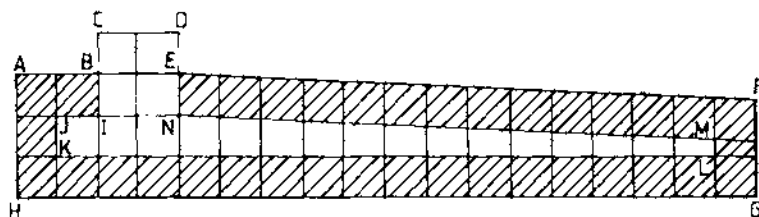


Figure 1.2.3 Example of metal solidification. *BIJKLMNE*: interface between metal and mould.

With the boundary, interface and initial conditions completely specified, the heat conduction equation can, in theory, be solved and the time variation of the temperature through the mould and the metal determined. The values of the thermal properties k and ρc of the metal and the mould will need to be specified before the solution process can begin.

1.3 VARIATIONAL FORMULATION

In the preceding section, we have outlined the classical formulation of the problem of heat conduction. This formulation is the starting point for certain approximate solution methods, such as the finite difference method, but is not appropriate for the finite element methods to be discussed in this text. For this reason, we will now replace the classical formulation by an equivalent variational formulation. It is this formulation which we shall employ as the starting point for the development of approximate solution techniques.

1.3.1 Steady problems

To illustrate the construction of a variational formulation, we begin by considering steady-state heat conduction in a domain, Ω , which consists of a single material. As in Figure 1.2.1 the temperature, T , will be assumed to be prescribed on the portion Γ_1 of the boundary curve, Γ , and the normal heat flux, q , will be prescribed on the remaining portion, Γ_2 . The classical statement of this problem is then: determine $T(\mathbf{x})$ such that

$$\text{div}(k \text{ grad } T) + Q = 0 \quad \text{in } \Omega \quad (1.3.1)$$

$$T = f(\mathbf{x}) \quad \text{for } \mathbf{x} \text{ on } \Gamma_1 \quad (1.3.2)$$

$$q(\mathbf{x}) = - (k \text{ grad } T) \cdot \mathbf{n} = -k \frac{\partial T}{\partial n} = \mathbf{N}(\mathbf{x}, T) \quad \text{for } \mathbf{x} \text{ on } \Gamma_2 \quad (1.3.3)$$

Here $k, f, \mathbf{N}$ and Q are prescribed functions. A variational formulation of the

problem requires the introduction of a set, $\mathcal{T}$, of trial functions and a set, $\mathcal{W}$, of weighting functions. At this point, we will assume that $\mathcal{T}$ consists of all functions that satisfy the boundary conditions of (1.3.2) and (1.3.3) and a strong variational formulation of the above problem is then: find T in the trial function set $\mathcal{T}$ such that

$$\int_{\Omega} [\operatorname{div}(k \operatorname{grad} T) + Q] W \, d\Omega = 0 \quad (1.3.4)$$

for each W in the weighting function set $\mathcal{W}$. An additional requirement now is that the functions which belong to the trial and weighting function sets must be sufficiently continuous so that the integral appearing in equation (1.3.4) is well defined. If we work with the largest possible sets, then the appropriate requirement here is that the weighting functions may be discontinuous but the trial functions must have continuous first derivatives. We use the notation that the weighting functions must show C^{-1} continuity and the trial functions must show C^1 continuity. The trial and weighting function sets can then be completely defined by

$$\mathcal{T} = \left\{ T \mid T = f(\mathbf{x}) \text{ on } \Gamma_1; q(\mathbf{x}) = -k \frac{\partial T}{\partial n} = \mathbf{N}(\mathbf{x}, T) \text{ on } \Gamma_2; T \text{ is } C^1 \right\} \quad (1.3.5)$$

$$\mathcal{W} = \{ W \mid W \text{ is } C^{-1} \} \quad (1.3.6)$$

By the fundamental theorem of the variational calculus, the function $T(\mathbf{x})$ which satisfies the formulation of equation (1.3.4) is a solution of (1.3.1). The definition of the trial function set automatically ensures that this function also satisfies the conditions (1.3.2) and (1.3.3). Hence, the function $T(\mathbf{x})$ which satisfies this variational formulation is also the function which is the solution of the classical statement. For this variational formulation, the boundary conditions of equations (1.3.2) and (1.3.3) are termed essential, as they have to be satisfied by all members of the trial function set.

In the formulation presented above, we have determined the required solution from among the set of trial functions which satisfies the problem boundary conditions and have used the variational statement to impose the additional requirement of the satisfaction of the governing differential equation. However, alternative variational formulations of the problem, in which we widen the definition of the trial function set, are also possible. For example, relaxing completely the requirement that members of the trial function set should satisfy the problem boundary conditions, we can work with the trial and weighting function sets

$$\mathcal{T} = \{ T \mid T \text{ is } C^1 \} \quad (1.3.7)$$

$$\mathcal{W} = \{ W \mid W \text{ is } C^{-1} \} \quad (1.3.8)$$

and the variational formulation: find T in $\mathcal{T}$ such that

$$\int_{\Omega} [\text{div}(k \text{ grad } T) + Q] W \, d\Omega + \Psi \int_{\Gamma_1} [T - f] W \, d\Gamma + \Phi \int_{\Gamma_2} \left[k \frac{\partial T}{\partial n} + \aleph \right] W \, d\Gamma = 0 \quad (1.3.9)$$

for every W in $\mathcal{W}$. In this equation, Ψ and Φ are arbitrary constants. The fundamental theorem of the variational calculus can again be used to deduce that the function $T(\mathbf{x})$, which is the solution of this formulation, will satisfy equation (1.3.1) everywhere in Ω , the boundary condition of equation (1.3.2) everywhere on Γ_1 , and the boundary condition of equation (1.3.3) everywhere on Γ_2 . Hence this function is also the solution of the classical formulation. Here, the boundary conditions of (1.3.2) and (1.3.3) can be termed natural, as they do not have to be satisfied by the members of the trial function set but are imposed via the variational formulation.

An alternative, so-called weak, formulation of the problem, results from manipulating equation (1.3.9), assuming that the functions involved are sufficiently smooth. Using the result

$$[\text{div}(k \text{ grad } T)] W = \text{div}[W k \text{ grad } T] - k \text{ grad } T \cdot \text{grad } W \quad (1.3.10)$$

and the Divergence Theorem, we produce, from equation (1.3.9), the equation

$$\begin{aligned} \int_{\Gamma_1 + \Gamma_2} k \frac{\partial T}{\partial n} W \, d\Gamma - \int_{\Omega} [k \text{ grad } T \cdot \text{grad } W - QW] \, d\Omega \\ + \Psi \int_{\Gamma_1} [T - f] W \, d\Gamma + \Phi \int_{\Gamma_2} \left[k \frac{\partial T}{\partial n} + \aleph \right] W \, d\Gamma = 0 \end{aligned} \quad (1.3.11)$$

which can be simplified, if we adopt the value $\Phi = -1$, to

$$\begin{aligned} \int_{\Omega} k \text{ grad } T \cdot \text{grad } W \, d\Omega - \int_{\Omega} QW \, d\Omega + \int_{\Gamma_1} W k \frac{\partial T}{\partial n} \, d\Gamma - \int_{\Gamma_2} W \aleph \, d\Gamma \\ + \Psi \int_{\Gamma_1} [T - f] W \, d\Gamma \end{aligned} \quad (1.3.12)$$

A further simplification is possible if we also require that the members of the trial function set, $\mathcal{T}$, satisfy the boundary condition of equation (1.3.2). In this case, the variational formulation becomes: find T in $\mathcal{T}$ such that

$$\int_{\Omega} k \text{ grad } T \cdot \text{grad } W \, d\Omega - \int_{\Omega} QW \, d\Omega + \int_{\Gamma_1} W k \frac{\partial T}{\partial n} \, d\Gamma - \int_{\Gamma_2} W \aleph \, d\Gamma \quad (1.3.13)$$

for every W in the weighting function set $\mathcal{W}$. It is apparent that, for this variational formulation, the continuity requirements which are placed upon the members of the trial function set are reduced, whereas increased continuity requirements are placed upon members of the weighting function set. In fact, the

requirement here is that members of both sets must be continuous (i.e. they must exhibit C^0 continuity) that the sets can be defined as

$$\mathcal{T} = \{T | T = f(\mathbf{x}) \text{ on } \Gamma_1; T \text{ is } C^0\} \quad (1.3.14)$$

$$\mathcal{W} = \{W | W \text{ is } C^0\} \quad (1.3.15)$$

A further simplification follows if the set of weighting functions is restricted to consist of those functions W which satisfy the homogeneous form of the boundary condition of equation (1.3.2), i.e.

$$W = 0 \quad \text{on } \Gamma_1 \quad (1.3.16)$$

Now, equation (1.3.13) implies that we look for T in $\mathcal{T}$ such that

$$\int_{\Omega} k \text{grad } T \cdot \text{grad } W \, d\Omega = \int_{\Omega} Q W \, d\Omega - \int_{\Gamma_2} W \mathbf{q} \cdot \mathbf{n} \, d\Gamma \quad (1.3.17)$$

for every W in $\mathcal{W}$. For this formulation, the appropriate trial and weighting function sets are

$$\mathcal{T} = \{T | T = f(\mathbf{x}) \text{ on } \Gamma_1; T \text{ is } C^0\} \quad (1.3.18)$$

$$\mathcal{W} = \{W | W = 0 \text{ on } \Gamma_1; W \text{ is } C^0\} \quad (1.3.19)$$

and it is apparent that there is now a close relationship between the two function sets. This close relationship is an attractive feature of this weak formulation for the heat conduction problem. This weak formulation is frequently preferred for both mathematical reasons and for practical applications.

For the weak formulation of equation (1.3.17)–(1.3.19), the boundary condition (1.3.2) is an essential boundary condition, as it is directly imposed upon the members of the trial function set. On the other hand, the boundary condition (1.3.3), which is automatically satisfied by the function $T(\mathbf{x})$ which is determined by the formulation, is a natural boundary condition as it is not directly imposed upon the members of the trial function set. The identification of any natural boundary conditions corresponding to a particular variational formulation can have important practical implications when an approximate solution of the problem is being attempted, as it can simplify considerably the construction of the trial function set.

1.3.2 Multi-material problems

For steady heat conduction in a region Ω consisting of two (or more) separate materials, it has been shown in equation (1.2.20) that the classical formulation requires the satisfaction of certain additional conditions at material interfaces. To consider the effect of this on our variational formulation, we assume that Ω is made up of two regions, Ω_1 and Ω_2 , with material interface $\tilde{\Gamma}$, as shown in Figure 1.2.2. The classical formulation is then: determine a continuous function

$T(\mathbf{x})$ such that

$$\operatorname{div}(k \operatorname{grad} T) + Q = 0 \quad \text{in } \Omega \quad (1.3.20)$$

$$T = f(\mathbf{x}) \quad \text{for } \mathbf{x} \text{ on } \Gamma_1 \quad (1.3.21)$$

$$q(\mathbf{x}) = -(k \operatorname{grad} T) \cdot \mathbf{n} = -k \frac{\partial T}{\partial n} = \mathbf{N}(\mathbf{x}, T) \quad \text{for } \mathbf{x} \text{ on } \Gamma_2 \quad (1.3.22)$$

$$(k \operatorname{grad} T)_1 \cdot \tilde{\mathbf{n}} = (k \operatorname{grad} T)_2 \cdot \tilde{\mathbf{n}} \quad \text{for } \mathbf{x} \text{ on } \tilde{\Gamma} \quad (1.3.23)$$

where the subscripts 1 and 2 refer to conditions appropriate to the regions Ω_1 and Ω_2 respectively and where $k_1, k_2, f, \mathbf{N}$ and Q are prescribed functions. We will demonstrate that an appropriate weak formulation for this problem is again the formulation implied by (1.3.17)–(1.3.19). This can be proved by noting that, using equation (1.3.10),

$$\int_{\Omega} k \operatorname{grad} T \cdot \operatorname{grad} W \, d\Omega = \int_{\Omega_1 \cup \Omega_2} [\operatorname{div}(W k \operatorname{grad} T) - W \operatorname{div}(k \operatorname{grad} T)] \, d\Omega \quad (1.3.24)$$

and it follows, using the Divergence Theorem over the regions Ω_1 and Ω_2 separately, that

$$\begin{aligned} \int_{\Omega} k \operatorname{grad} T \cdot \operatorname{grad} W \, d\Omega &= - \int_{\Omega} W \operatorname{div}(k \operatorname{grad} T) \, d\Omega + \int_{\Gamma_2} W k \operatorname{grad} T \cdot \mathbf{n} \, d\Gamma \\ &\quad + \int_{\tilde{\Gamma}} W [k \operatorname{grad} T \cdot \tilde{\mathbf{n}}] \, d\Gamma \end{aligned} \quad (1.3.25)$$

where, for $\mathbf{x}$ on $\tilde{\Gamma}$, $[f(\mathbf{x})] = f_2(\mathbf{x}) - f_1(\mathbf{x})$ denotes the change in f across $\tilde{\Gamma}$. Thus the variational formulation of equations (1.3.17)–(1.3.19) is, in this case, equivalent to finding the function T in the trial function set $\mathcal{T}$ of equation (1.3.18) which is such that

$$\begin{aligned} &- \int_{\Omega} [\operatorname{div}(k \operatorname{grad} T) + Q] W \, d\Omega + \int_{\Gamma_2} [k \operatorname{grad} T \cdot \mathbf{n} + \mathbf{N}] W \, d\Gamma \\ &\quad + \int_{\tilde{\Gamma}} [k \operatorname{grad} T \cdot \tilde{\mathbf{n}}] W \, d\Gamma = 0 \end{aligned} \quad (1.3.26)$$

for all W in the weighting function set $\mathcal{W}$ of equation (1.3.19). This statement implies that the function $T(\mathbf{x})$ which satisfies the weak variational formulation will also be the solution of the classical problem defined by equations (1.3.20)–(1.3.23). The important feature to note is that both the boundary condition on Γ_2 and the interface condition on $\tilde{\Gamma}$ are natural boundary conditions for this weak formulation.

1.3.3 Transient problems

The inclusion of the effects of the time variable can be made directly within the variational formulations outlined above, provided that Ω and Γ now refer to the full space-time domain and its boundaries respectively. However, this is not the approach generally followed in this text, where the main intention is to treat the space and the time domains separately. In this case, the time t is considered as a real parameter and we develop a family of one-parameter variational formulations in t . We illustrate this approach by considering a problem, which is to be solved for a single material within a spatial domain Ω with spatial boundary Γ , whose classical formulation is: find $T(\mathbf{x}, t)$ such that

$$\operatorname{div}(k \operatorname{grad} T) + Q = \rho c \frac{\partial T}{\partial t} \quad \text{for } \mathbf{x} \text{ in } \Omega, t > 0 \quad (1.3.27)$$

$$T = f(\mathbf{x}, t) \quad \text{for } \mathbf{x} \text{ on } \Gamma_1, t > 0 \quad (1.3.28)$$

$$q(\mathbf{x}, t) = -(k \operatorname{grad} T) \cdot \mathbf{n} = -k \frac{\partial T}{\partial n} = \mathbf{N}(\mathbf{x}, T, t) \quad \text{for } \mathbf{x} \text{ on } \Gamma_2, t > 0 \quad (1.3.29)$$

$$T = g(\mathbf{x}) \quad \text{for } \mathbf{x} \text{ in } \Omega, t = 0 \quad (1.3.30)$$

Here $\rho, c, k, f, g, \mathbf{N}$ and Q are prescribed functions. To determine the solution over the interval $0 < t < \tau$, we define trial and weighting function sets by

$$\mathcal{T} = \{T(\mathbf{x}, t) | T = f(\mathbf{x}, t) \text{ on } \Gamma_1, 0 < t < \tau; T = g(\mathbf{x}) \text{ in } \Omega, t = 0; T \text{ is } C^0\} \quad (1.3.31)$$

$$\mathcal{W} = \{W(\mathbf{x}) | W = 0 \text{ on } \Gamma_1; W \text{ is } C^0\} \quad (1.3.32)$$

and we consider the variational formulation: find $T(\mathbf{x}, t)$ in $\mathcal{T}$ such that

$$\int_{\Omega} \left[\rho c \frac{\partial T}{\partial t} W + k \operatorname{grad} T \cdot \operatorname{grad} W \right] d\Omega = \int_{\Omega} Q W d\Omega - \int_{\Gamma_2} W \mathbf{N} d\Gamma \quad (1.3.33)$$

for every t in the range $0 < t < \tau$ and for every $W(\mathbf{x})$ in $\mathcal{W}$. By employing equation (1.3.24) and the Divergence Theorem on the left-hand side of this equation, it can be demonstrated that the solution T of this variational problem will also be the solution of the classical problem of equations (1.3.27)–(1.3.30). It should be noted that the boundary condition of equation (1.3.28) and the initial condition of equation (1.3.30) are essential conditions within this formulation, while the boundary condition of equation (1.3.29) is natural.

1.4 THE GALERKIN APPROXIMATE SOLUTION METHOD

The integral weak variational formulation of the heat conduction problem is the starting point for the development of the Galerkin method for constructing

approximate solutions. We illustrate the approach by concentrating upon a particular example and we consider the solution of the steady linear heat conduction problem of equations (1.3.1)–(1.3.3), using initially the variational formulation of equation (1.3.17)–(1.3.19). We begin by determining any function $\psi(\mathbf{x})$ which satisfies the essential boundary condition on Γ_1 i.e.

$$\psi(\mathbf{x}) = f(\mathbf{x}) \quad \text{for } \mathbf{x} \text{ on } \Gamma_1 \quad (1.4.1)$$

and choose a complete set of continuous approximating functions $N_1, N_2, \dots$, such that each member N_j of this set satisfies the homogeneous form of the boundary condition on Γ_1 i.e.

$$N_j(\mathbf{x}) = 0 \quad \text{for } \mathbf{x} \text{ on } \Gamma_1 \quad (1.4.2)$$

The completeness requirement guarantees that these approximating functions form a basis for the trial and weighting function sets of (1.3.18) and (1.3.19), i.e.

$$\mathcal{T} = \left\{ T \mid T = \psi + \sum_{j=1}^{\infty} a_j N_j; T = f(\mathbf{x}) \text{ on } \Gamma_1; T \text{ is } C^0 \right\} \quad (1.4.3)$$

$$\mathcal{W} = \left\{ W \mid W = \sum_{j=1}^{\infty} b_j N_j; W = 0 \text{ on } \Gamma_1; W \text{ is } C^0 \right\} \quad (1.4.4)$$

where $a_1, a_2, \dots$, and $b_1, b_2, \dots$, are constants. To construct an approximate solution to the problem, we consider subsets $\mathcal{T}^{(M)}$ and $\mathcal{W}^{(M)}$, of finite dimension M , of the trial and weighting function sets, which are defined by

$$\mathcal{T}^{(M)} = \left\{ \hat{T} \mid \hat{T} = \psi + \sum_{j=1}^M a_j N_j; \hat{T} = f(\mathbf{x}) \text{ on } \Gamma_1; \hat{T} \text{ is } C^0 \right\} \quad (1.4.5)$$

$$\mathcal{W}^{(M)} = \left\{ W \mid W = \sum_{j=1}^M b_j N_j; W = 0 \text{ on } \Gamma_1; W \text{ is } C^0 \right\} \quad (1.4.6)$$

The Galerkin approximate solution is determined by employing the weak variational formulation of (1.3.17)–(1.3.19) as follows: find $\hat{T}$ in $\mathcal{T}^{(M)}$ such that

$$\int_{\Omega} k \text{grad } \hat{T} \cdot \text{grad } W \, d\Omega - \int_{\Omega} Q W \, d\Omega - \int_{\Gamma_2} W \mathcal{N} \, d\Gamma \quad (1.4.7)$$

for every W in the weighting function set $\mathcal{W}^{(M)}$. It can be observed from (1.4.6) that every W in $\mathcal{W}^{(M)}$ is formed as a linear combination of the approximating functions $N_1, N_2, \dots, N_M$ and hence the requirement of equation (1.4.7) is equivalent to the requirement that the function $\hat{T}$ be such that

$$\int_{\Omega} k \text{grad } \hat{T} \cdot \text{grad } N_i \, d\Omega - \int_{\Omega} Q N_i \, d\Omega - \int_{\Gamma_2} N_i \mathcal{N} \, d\Gamma = 0 \quad \text{for } i = 1, 2, \dots, M \quad (1.4.8)$$

Inserting the assumed form for $\hat{T}$ from equation (1.4.5), it follows that

$$\begin{aligned} \int_{\Omega} k \operatorname{grad} \psi \cdot \operatorname{grad} N_i \, d\Omega + \sum_{j=1}^M \left\{ \int_{\Omega} k \operatorname{grad} N_j \cdot \operatorname{grad} N_i \, d\Omega \right\} a_j \\ = \int_{\Omega} Q N_i \, d\Omega - \int_{\Gamma_2} N_i \bar{N} \, d\Gamma \quad \text{for } i = 1, 2, \dots, M \end{aligned} \quad (1.4.9)$$

This set of M equations can be expressed in the matrix form

$$\mathbf{K} \mathbf{a} = \mathbf{r} \quad (1.4.10)$$

where $\mathbf{a}$ is the vector defined by

$$\mathbf{a}^T = (a_1, a_2, \dots, a_M) \quad (1.4.11)$$

and $\mathbf{r}$ is a vector with typical entry

$$r_i = \int_{\Omega} Q N_i \, d\Omega - \int_{\Gamma_2} N_i \bar{N} \, d\Gamma - \int_{\Omega} k \operatorname{grad} \psi \cdot \operatorname{grad} N_i \, d\Omega \quad (1.4.12)$$

The $M \times M$ matrix $\mathbf{K}$ will, in general, be full with typical entry

$$K_{ij} = \int_{\Omega} k \operatorname{grad} N_j \cdot \operatorname{grad} N_i \, d\Omega \quad (1.4.13)$$

The solution of the equation system (1.4.10) for the constants $a_1, a_2, \dots, a_M$ completes the determination of the Galerkin approximate solution $\hat{T}$ in the subset $\mathcal{T}^{(M)}$.

1.4.1 Optimality of the Galerkin approximation

For the problem of steady linear heat conduction, we are able to demonstrate, rather easily, that the Galerkin approximation possesses a certain optimality property. If

$$E = E(\mathbf{x}) = T - \hat{T} \quad (1.4.14)$$

denotes the error in the Galerkin approximation and

$$E_S = E_S(\mathbf{x}) = T - \hat{S} \quad (1.4.15)$$

is the error which results when any other function $\hat{S}$ in $\mathcal{T}^{(M)}$ is used as an approximation to T , then the optimality result can be expressed as

$$\|E\| \leq \|E_S\| \quad (1.4.16)$$

where $\|\cdot\|$ denotes an appropriate measure, i.e. among the whole set of functions $\mathcal{T}^{(M)}$, the function constructed according to the Galerkin procedure is the best approximation possible according to the measure.

From equation (1.3.17), the exact solution T satisfies

$$\int_{\Omega} k \operatorname{grad} T \cdot \operatorname{grad} W \, d\Omega = \int_{\Omega} QW \, d\Omega - \int_{\Gamma_2} WN \, d\Gamma \quad (1.4.17)$$

for every W in the weighting function set $\mathcal{W}$. Since $\mathcal{W}^{(M)}$ is a subset of $\mathcal{W}$, this equation must also be valid for every W in $\mathcal{W}^{(M)}$. From equation (1.4.7), the Galerkin approximation $\hat{T}$ satisfies

$$\int_{\Omega} k \operatorname{grad} \hat{T} \cdot \operatorname{grad} W \, d\Omega = \int_{\Omega} QW \, d\Omega - \int_{\Gamma_2} WN \, d\Gamma \quad (1.4.18)$$

for every W in the weighting function set $\mathcal{W}^{(M)}$. Subtracting these two equations, it follows that

$$\int_{\Omega} k \operatorname{grad} E \cdot \operatorname{grad} W \, d\Omega = 0 \quad (1.4.19)$$

for every W in the weighting function set $\mathcal{W}^{(M)}$. If we define a measure according to

$$\|f\|^2 = \int_{\Omega} k \operatorname{grad} f \cdot \operatorname{grad} f \, d\Omega \quad (1.4.20)$$

the result

$$E_S = T - \hat{S} = (T - \hat{T}) + (\hat{T} - \hat{S}) = E + (\hat{T} - \hat{S}) \quad (1.4.21)$$

leads to

$$\begin{aligned} \|E_S\|^2 &= \int_{\Omega} k \operatorname{grad} E \cdot \operatorname{grad} E \, d\Omega + 2 \int_{\Omega} k \operatorname{grad} E \cdot \operatorname{grad} (\hat{T} - \hat{S}) \, d\Omega \\ &\quad + \int_{\Omega} k \operatorname{grad} (\hat{T} - \hat{S}) \cdot \operatorname{grad} (\hat{T} - \hat{S}) \, d\Omega \end{aligned} \quad (1.4.22)$$

Since both $\hat{T}$ and $\hat{S}$ belong to $\mathcal{T}^{(M)}$, it can be observed, from equations (1.4.5) and (1.4.6), that $\hat{T} - \hat{S}$ belongs to $\mathcal{W}^{(M)}$. Hence, from equation (1.4.19), we deduce that

$$\int_{\Omega} k \operatorname{grad} E \cdot \operatorname{grad} (\hat{T} - \hat{S}) \, d\Omega = 0 \quad (1.4.23)$$

and it follows that equation (1.4.22) can be written as

$$\|E_S\|^2 = \|E\|^2 + \|\hat{T} - \hat{S}\|^2 \quad (1.4.24)$$

The optimality result of equation (1.4.16) is now apparent since, from the definition of the measure in equation (1.4.20), $\|\hat{T} - \hat{S}\|^2$ is strictly non-negative.

1.4.2 Convergence of the Galerkin approximation

The Galerkin method, as outlined above, can be used with confidence in practice to construct approximate solutions to problems involving heat conduction if the convergence of the procedure can be demonstrated, i.e. if it can be proved that $\hat{T} \rightarrow T$ as the dimension M of the subsets $\mathcal{F}^{(M)}$ and $\mathcal{H}^{(M)}$ increases. The proof of convergence is not straightforward, but it is possible to show that, for linear heat conduction problems, convergence occurs according to the measure defined by equation (1.4.20), i.e. that $\|T - \hat{T}\| \rightarrow 0$ as $M \rightarrow \infty$.

1.5 FINITE ELEMENT APPROXIMATING FUNCTIONS IN ONE DIMENSION

It is apparent from the preceding discussion that the analyst who is interested in obtaining approximate solutions to heat conduction problems by the variational approach is immediately faced with the difficulty of choosing a suitable set of approximating functions. The finite element method helps the analyst in this respect by providing a systematic way of constructing a polynomial approximating function set of any desired dimension, M , by utilising the results of interpolation theory.

Consider the construction of an approximate solution to the one-dimensional steady linear heat conduction problem defined by

$$\frac{d}{d\xi} \left(k \frac{dT}{d\xi} \right) + Q = 0 \quad -1 < \xi < 1 \quad (1.5.1)$$

$$T = f \quad \text{at } \xi = -1 \quad (1.5.2)$$

$$q = -k \frac{dT}{dn} = -k \frac{dT}{d\xi} = \aleph \quad \text{at } \xi = 1 \quad (1.5.3)$$

where k and Q are prescribed functions and f and $\aleph$ are prescribed constants. To obtain a solution by the finite element method, we begin by placing a set of $M + 1$ points, or nodes, at selected points $\xi_1, \xi_2, \dots, \xi_{M+1}$ on the region, or element, defined by $-1 \leq \xi \leq 1$. One of these nodes is located at $\xi = -1$ and another is located at $\xi = 1$. For present purposes, the nodes are numbered according to the numbering system shown in Figure 1.5.1. With each node j we define the associated Lagrange interpolation polynomial $\mathcal{L}_j^{(M)}(\xi)$, of degree M , according to the requirements that

$$\mathcal{L}_j^{(M)}(\xi) = \begin{cases} 1 & \text{if } \xi = \xi_j \\ 0 & \text{if } \xi = \xi_k, k \neq j \end{cases} \quad (1.5.4)$$

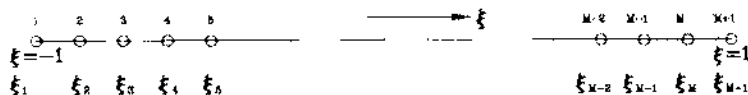


Figure 1.5.1 Node numbering for one-dimensional problem using linear elements.

which leads to the expression

$$\mathcal{L}_j^{(M)}(\xi) = \prod_{\substack{l=1 \\ l \neq j}}^{M+1} \frac{(\xi_l - \xi)}{(\xi_j - \xi_l)} \quad (1.5.5)$$

If we now adopt the choice $N_j = \mathcal{L}_j^{(M)}$, then the function

$$\hat{T} = \sum_{j=1}^{M+1} a_j N_j \quad (1.5.6)$$

is such that $\hat{T} = a_j$ when $\xi = \xi_j$, i.e. the coefficient a_j is the value of the function $\hat{T}$ at the point $\xi = \xi_j$. To denote this, we replace the representation of equation (1.5.6) by

$$\hat{T} = \sum_{j=1}^{M+1} T_j N_j \quad (1.5.7)$$

and adopt the definitions

$$\mathcal{T}^{(M)} = \left\{ \hat{T} \mid \hat{T} = \psi + \sum_{j=2}^{(M+1)} T_j N_j; \psi = f N_1; \hat{T} = f \text{ at } \xi = -1 \right\} \quad (1.5.8)$$

for the trial function set and

$$\mathcal{W}^{(M)} = \left\{ W \mid W = \sum_{j=2}^{(M+1)} b_j N_j; W = 0 \text{ at } \xi = -1 \right\} \quad (1.5.9)$$

for the weighting function set. Following the procedure outlined in Section 1.4, these finite dimensional trial and weighting functions sets can be used to determine an approximate solution to the problem of equations (1.5.1)–(1.5.3) by the Galerkin method. The resulting solution is the Galerkin finite element approximation which is obtained when a polynomial of order (or degree) M is employed over a single $M + 1$ noded finite element.

In the current context, the function N_j is termed the finite element shape function associated with node j and it is informative at this stage to examine some low order one-dimensional elements in greater detail.

1.5.1 The linear element

The linear element, for which $M = 1$, has two nodes, located at $\xi = -1$ and at $\xi = 1$, as shown in Figure 1.5.2. The corresponding shape functions are linear

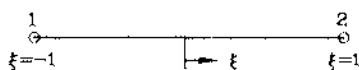


Figure 1.5.2 One-dimensional linear element.

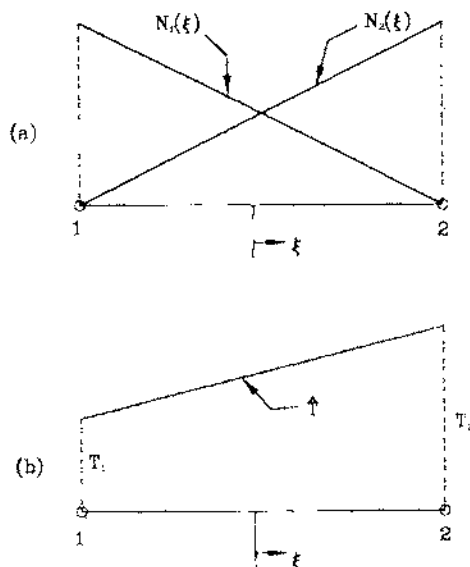


Figure 1.5.3 Shape functions and form of function for one-dimensional linear element.

over the element and equation (1.5.5) can be used to construct expressions for these functions as

$$N_1(\xi) = (1 - \xi)/2 \quad N_2(\xi) = (1 + \xi)/2 \quad (1.5.10)$$

These shape functions are shown in Figure 1.5.3, which also shows the form of the function $\hat{T}$ when nodal values T_1 and T_2 are prescribed.

The problem of equations (1.5.1)–(1.5.3) is defined over the region $-1 \leq \xi \leq 1$. Problems involving the more general region $x_1 \leq x \leq x_2$ can be handled by mapping the region $-1 \leq \xi \leq 1$ into the region of interest. An isoparametric mapping accomplishes this by employing the element shape functions and defining the mapping by

$$x = \sum_{j=1}^2 x_j N_j(\xi) \quad -1 \leq \xi \leq 1 \quad (1.5.11)$$

In this case, the trial and weighting function sets of (1.5.8) and (1.5.9) will become

$$\mathcal{T}^{(1)} = \{ \hat{T}(x); \hat{T}(x) = \psi + T_2 N_2(\xi); \psi = f N_1(\xi); \hat{T} = f \text{ at } x = x_1 \} \quad (1.5.12)$$

$$\mathcal{W}^{(1)} = \{ W(x); W(x) = b_2 N_2(\xi); W = 0 \text{ at } x = x_1 \} \quad (1.5.13)$$

with the relation between the coordinates x and ξ being given by the mapping of equation (1.5.11).

1.5.2 The quadratic element

When $M = 2$, the nodal shape functions are quadratic over the element. We consider the case where the three nodes are equally spaced and located at the points $\xi = -1$, $\xi = 0$ and $\xi = 1$. Using equation (1.5.5), the corresponding nodal shape functions are readily shown to be

$$N_1(\xi) = -\xi(1 - \xi)/2 \quad N_2(\xi) = (1 - \xi^2) \quad N_3(\xi) = \xi(1 + \xi)/2 \quad (1.5.14)$$

These shape functions are displayed in Figure 1.5.4, which also shows the form of the function $\hat{T}$ when nodal values T_1 , T_2 and T_3 are prescribed.

Approximate solutions to problems involving the general region $x_1 \leq x \leq x_3$, with nodes located at x_1 , x_2 and x_3 , can also be handled by using a single quadratic element. This is accomplished by again mapping the region $-1 \leq \xi \leq 1$ into the region of interest and this can be achieved by the isoparametric mapping

$$x = \sum_{j=1}^3 x_j N_j(\xi) \quad -1 \leq \xi \leq 1 \quad (1.5.15)$$

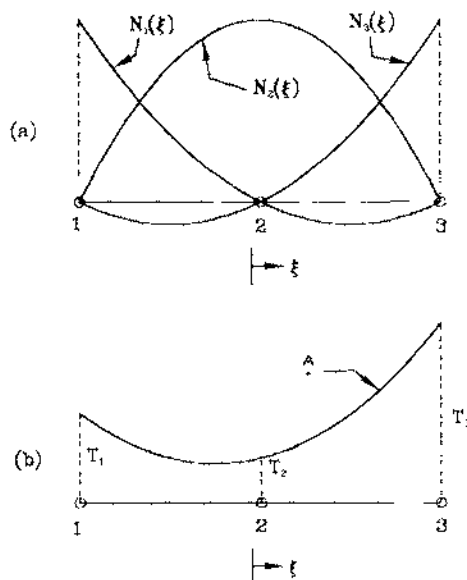


Figure 1.5.4 Shape functions and form of function for one-dimensional quadratic element.

The trial and weighting function sets of (1.5.8) and (1.5.9) now become

$$\mathcal{T}^{(2)} = \left\{ \hat{T}(x); \hat{T}(x) = \psi + \sum_{j=2}^3 T_j N_j(\xi); \psi = f N_1(\xi); \hat{T} = f \text{ at } x = x_1 \right\} \quad (1.5.16)$$

$$\mathcal{W}^{(2)} = \left\{ W(x); W(x) = \sum_{j=2}^3 b_j N_j(\xi); W = 0 \text{ at } x = x_1 \right\} \quad (1.5.17)$$

with the relation between the coordinates x and ξ being given by the mapping of equation (1.5.15).

1.6 FINITE ELEMENT APPROXIMATING FUNCTIONS IN TWO DIMENSIONS

Now consider the construction of an approximate solution to a two-dimensional heat conduction problem on a square region defined by

$$\frac{\partial}{\partial \xi} \left(k \frac{\partial T}{\partial \xi} \right) + \frac{\partial}{\partial \eta} \left(k \frac{\partial T}{\partial \eta} \right) + Q = 0 \quad -1 < \xi, \eta < 1 \quad (1.6.1)$$

$$T = f(\xi) \quad \text{on } \eta = -1 \quad (1.6.2)$$

$$q = -k \frac{\partial T}{\partial n} = \mathcal{N} \quad \text{on } \xi = \pm 1 \text{ and on } \eta = 1 \quad (1.6.3)$$

Here $k, f, \mathcal{N}$ and Q are prescribed functions. If an approximate solution to this problem is to be sought by applying the Galerkin procedure, a suitable polynomial approximating function set can be obtained by direct extension of the one-dimensional finite element ideas introduced previously. We select a set of $M+1$ points $\xi_1, \xi_2, \dots, \xi_{M+1}$, such that $\xi_1 = -1$ and $\xi_{M+1} = 1$, and a set of $M+1$ points $\eta_1, \eta_2, \dots, \eta_{M+1}$, such that $\eta_1 = -1$ and $\eta_{M+1} = 1$. Nodes, numbered so that the first $4M$ lie on the boundary and the first $M+1$ lie on the line $\eta = -1$, as shown in Figure 1.6.1, are located at the points (ξ_k, η_l) and the associated Lagrange interpolation polynomials $\mathcal{L}_k^{(M)}(\xi)$ and $\mathcal{L}_l^{(M)}(\eta)$ are constructed. For node j , located at (ξ_k, η_l) , the finite element shape function is defined by

$$N_j(\xi, \eta) = \mathcal{L}_k^{(M)}(\xi) \mathcal{L}_l^{(M)}(\eta) \quad (1.6.4)$$

It is readily demonstrated that this function has the value unity at node j and the value zero at all the other nodes. The trial function set for these approximating functions becomes

$$\mathcal{T}^{(M(M+1))} = \left\{ \hat{T}; \hat{T} = \psi + \sum_{j=M+2}^{(M+1)^2} T_j N_j; \psi = \sum_{j=1}^{M+1} f_j N_j; \hat{T} = \hat{f} \text{ on } \eta = -1 \right\} \quad (1.6.5)$$

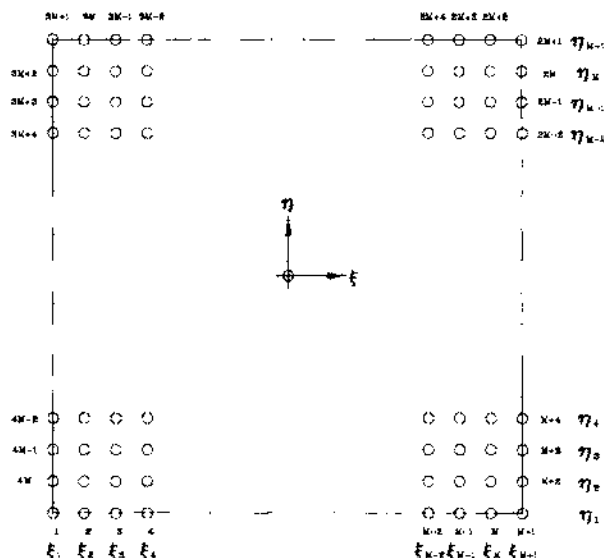


Figure 1.6.1 Node numbering for two-dimensional problem using linear elements.

and the corresponding weighting function set is

$$\mathcal{W}^{(M(M+1))} = \left\{ W \mid W = \sum_{j=M+2}^{(M+1)^2} b_j N_j; W = 0 \text{ on } \eta = -1 \right\} \quad (1.6.6)$$

Note now that the members $\hat{T}$ of the trial function set satisfy $\hat{T} = \hat{f}$ on the boundary $\eta = -1$, where

$$\hat{f}(\xi) = \sum_{j=1}^{M+1} f_j N_j(\xi, -1) \quad (1.6.7)$$

In this representation, f_j denotes the value of $f(\xi)$ at node j , i.e. at $\xi = \xi_j$, and (1.6.7) is the Lagrange interpolating polynomial of degree M which passes through the prescribed values $f_1, f_2, \dots, f_{M+1}$ at the points $\xi_1, \xi_2, \dots, \xi_{M+1}$.

It is informative at this stage to derive the shape functions and the corresponding trial and weighting function sets for some low order two-dimensional elements.

1.6.1 The linear element

The choice $M = 1$ leads to a four-noded square element with the nodes located at the corners, as shown in Figure 1.6.2. Equations (1.5.5) and (1.6.4) can be used to show that the corresponding nodal shape functions are

$$\begin{aligned} N_1 &= (1 - \xi)(1 - \eta)/4 & N_2 &= (1 + \xi)(1 - \eta)/4 \\ N_3 &= (1 + \xi)(1 + \eta)/4 & N_4 &= (1 - \xi)(1 + \eta)/4 \end{aligned} \quad (1.6.8)$$

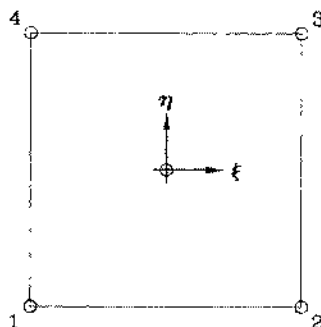


Figure 1.6.2 Four-noded square element.

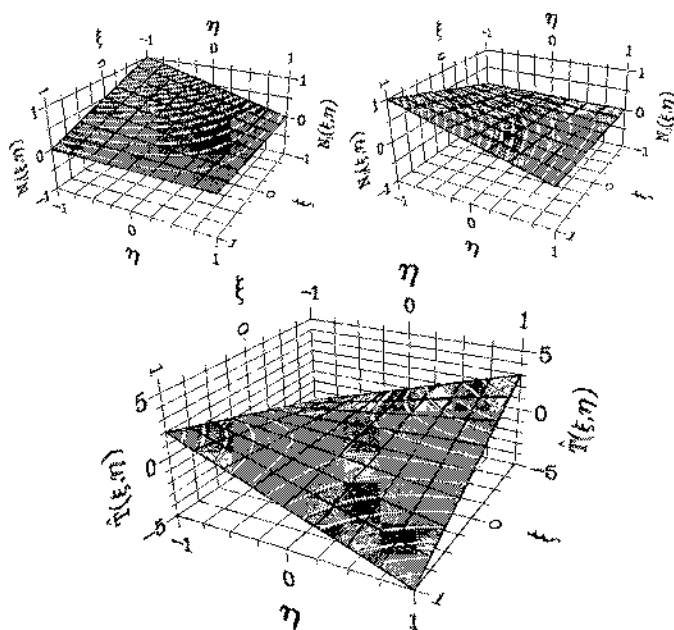


Figure 1.6.3 Form of shape functions for four-noded element.

The form of these shape functions is illustrated in Figure 1.6.3, which also shows the form of the function $\hat{T}$ when nodal values T_1, T_2, T_3 and T_4 are prescribed.

For the problem

$$\frac{\partial}{\partial x} \left(k \frac{\partial T}{\partial x} \right) + \frac{\partial}{\partial y} \left(k \frac{\partial T}{\partial y} \right) - Q = 0 \quad \mathbf{x} \text{ in } \Omega \quad (1.6.9)$$

$$T = f(\mathbf{x}) \quad \mathbf{x} \text{ on } \Gamma_1 \quad (1.6.10)$$

$$q = -k \frac{\partial T}{\partial n} = \mathbf{N}(\mathbf{x}) \quad \mathbf{x} \text{ on } \Gamma_2 \quad (1.6.11)$$

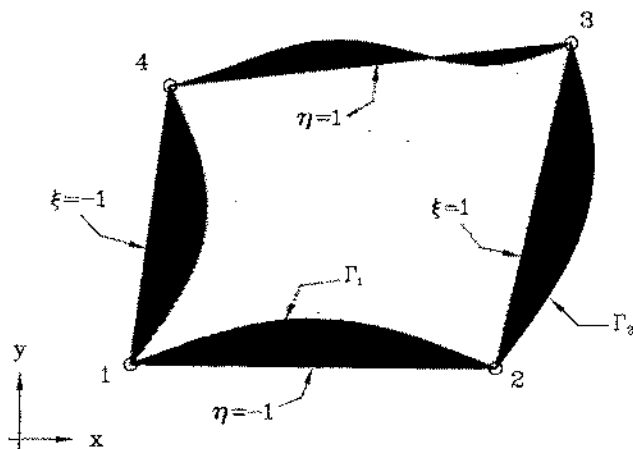


Figure 1.6.4 Approximation of a given region by a four-noded quadrilateral.

the region Ω is approximated by a single four-noded, straight-sided element. To accomplish the approximation we first select four points (x_j, y_j) ; $j = 1, 2, 3, 4$, on the boundary Γ , with (x_1, y_1) and (x_2, y_2) located at the end points of Γ_1 . The approximation to Ω is then the region formed by the isoparametric mapping of the region $-1 \leq \xi, \eta \leq 1$ to the (x, y) plane, defined by

$$x = \sum_{j=1}^4 x_j N_j(\xi, \eta) \quad y = \sum_{j=1}^4 y_j N_j(\xi, \eta) \quad (1.6.12)$$

This process is illustrated in Figure 1.6.4, with the darkly shaded area indicating the error which is made in approximating a region in this manner.

When a single linear element is used in this way to construct a Galerkin approximate solution to problem (1.6.9)–(1.6.11), the trial and weighting function sets are

$$\mathcal{T}^{(2)} = \left\{ \hat{T}(x, y) \mid \hat{T}(x, y) = \psi + \sum_{j=3}^4 T_j N_j(\xi, \eta); \psi = \sum_{j=1}^2 f_j N_j(\xi, \eta); \hat{T} = \hat{f} \text{ on } \eta = -1 \right\} \quad (1.6.13)$$

$$\mathcal{W}^{(2)} = \left\{ W(x, y) \mid W(x, y) = \sum_{j=3}^4 b_j N_j(\xi, \eta); W = 0 \text{ on } \eta = -1 \right\} \quad (1.6.14)$$

with the relation between the coordinates (x, y) and (ξ, η) being defined by the mapping of equation (1.6.12). In this way, the Galerkin solution is constructed over the approximated region in the (x, y) plane.

1.6.2 The quadratic element

We consider the choice $M = 2$, with the nodes located at the corners of the element, at the mid-points of the element sides and at the element centre. The nine nodes are numbered as in Figure 1.6.5, and the corresponding nodal shape functions follow from equations (1.5.5) and (1.6.4) as

$$\begin{aligned} N_1 &= (\xi^2 - \xi)(\eta^2 - \eta)/4 & N_2 &= (1 - \xi^2)(\eta^2 - \eta)/2 \\ N_3 &= (\xi^2 + \xi)(\eta^2 - \eta)/4 & N_4 &= (\xi^2 + \xi)(1 - \eta^2)/2 \\ N_5 &= (\xi^2 + \xi)(\eta^2 + \eta)/4 & N_6 &= (1 - \xi^2)(\eta^2 + \eta)/2 \\ N_7 &= (\xi^2 - \xi)(\eta^2 + \eta)/4 & N_8 &= (\xi^2 - \xi)(1 - \eta^2)/2 \\ N_9 &= (1 - \xi^2)(1 - \eta^2) \end{aligned} \quad (1.6.15)$$

The form of these shape functions is illustrated in Figure 1.6.6, which also shows the form of the function $\hat{T}$ when nodal values $T_1, T_2, \dots, T_9$ are prescribed.

An approximate solution to the problem of equations (1.6.9)–(1.6.11) may be constructed by first employing an isoparametric mapping of a single nine-noded element to approximate the region of interest, Ω , in the (x, y) plane. This is accomplished by selecting eight points (x_j, y_j) , $j = 1, 2, \dots, 8$, on the boundary Γ , with (x_1, y_1) and (x_3, y_3) located at the end points of Γ_1 and (x_2, y_2) also located on Γ_1 . A point (x_9, y_9) is also selected on the interior of Ω . The isoparametric relation

$$x = \sum_{j=1}^9 x_j N_j(\xi, \eta) \quad y = \sum_{j=1}^9 y_j N_j(\xi, \eta) \quad (1.6.16)$$

is then used to map the region $-1 \leq \xi, \eta \leq 1$ into the (x, y) plane. This process is illustrated in Figure 1.6.7, with the darkly shaded area indicating the error which is made in approximating the region in this manner.

When a single quadratic element is used in this way to construct a Galerkin approximate solution to problem (1.6.9)–(1.6.11), the trial and weighting function

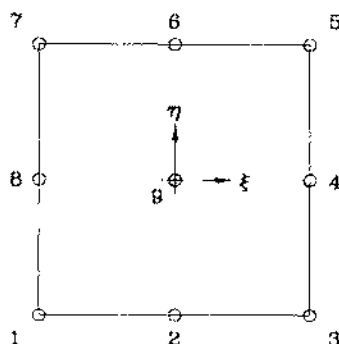


Figure 1.6.5 Nine-noded quadriilateral element.

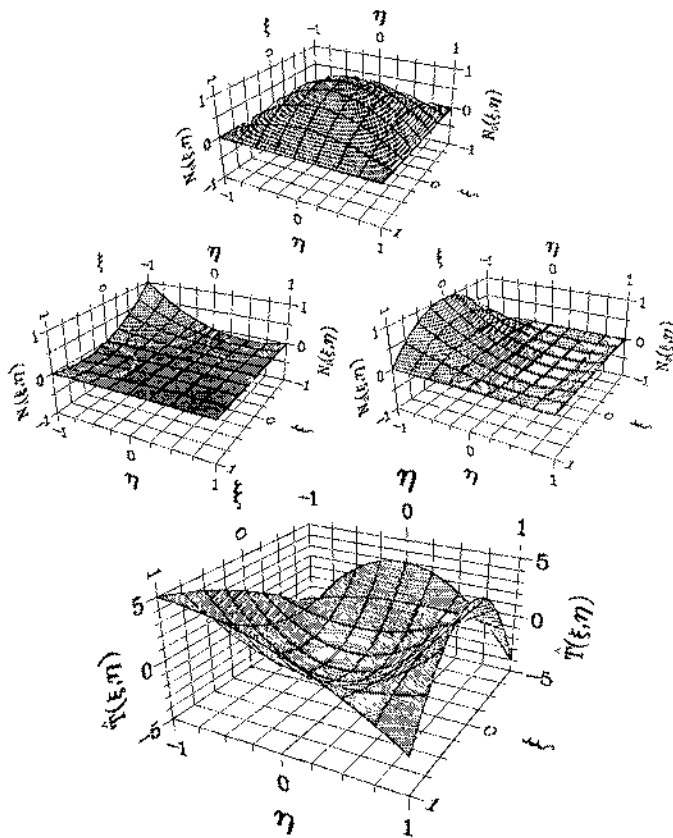


Figure 1.6.6 Form of shape functions for nine-noded quadrilateral element.

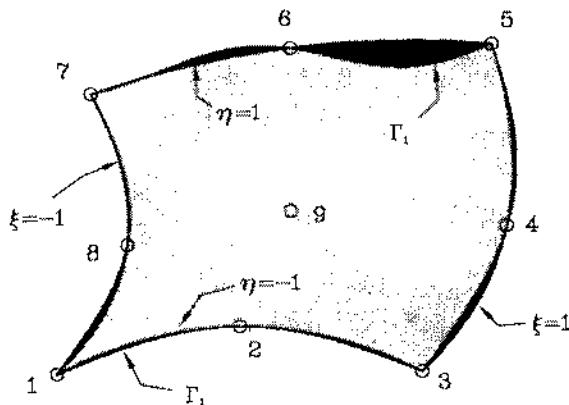


Figure 1.6.7 Approximation of a given region by a nine-noded quadrilateral.

sets are

$$\mathcal{F}^{(6)} = \left\{ \hat{T}(x, y) | \hat{T}(x, y) = \psi + \sum_{j=4}^9 T_j N_j(\xi, \eta); \psi = \sum_{j=1}^3 f_j N_j(\xi, \eta); \hat{T} = \hat{f} \text{ on } \eta = -1 \right\} \quad (1.6.17)$$

$$\mathcal{W}^{(6)} = \left\{ W(x, y) | W(x, y) = \sum_{j=4}^9 b_j N_j(\xi, \eta); W = 0 \text{ on } \eta = -1 \right\} \quad (1.6.18)$$

with the relation between the coordinates (x, y) and (ξ, η) being defined by the mapping of equation (1.6.16).

1.6.3 The quadratic serendipity element

The serendipity family of elements provides an alternative approach to the problem of defining polynomial approximating functions over a square region. These elements possess a reduced number of internal nodes and the shape function associated with node j is determined directly from the requirement that

$$N_j = \begin{cases} 1 & \text{at node } j \\ 0 & \text{at node } k; k \neq j \end{cases} \quad (1.6.19)$$

We illustrate this approach for an eight-noded element, with the nodes located as shown in Figure 1.6.8. In this case, it is easy to confirm that the corresponding nodal shape functions are

$$\begin{aligned} N_1 &= (1 - \xi)(1 - \eta)(-1 - \xi - \eta)/4 & N_2 &= (1 - \xi^2)(1 - \eta)/2 \\ N_3 &= (1 + \xi)(1 - \eta)(-1 + \xi - \eta)/4 & N_4 &= (1 + \xi)(1 - \eta^2)/2 \\ N_5 &= (1 + \xi)(1 + \eta)(-1 + \xi + \eta)/4 & N_6 &= (1 - \xi^2)(1 + \eta)/2 \\ N_7 &= (1 - \xi)(1 + \eta)(-1 - \xi + \eta)/4 & N_8 &= (1 - \xi)(1 + \eta^2)/2 \end{aligned} \quad (1.6.20)$$

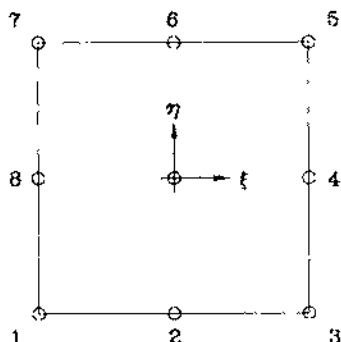


Figure 1.6.8 Eight-noded quadratic serendipity element.

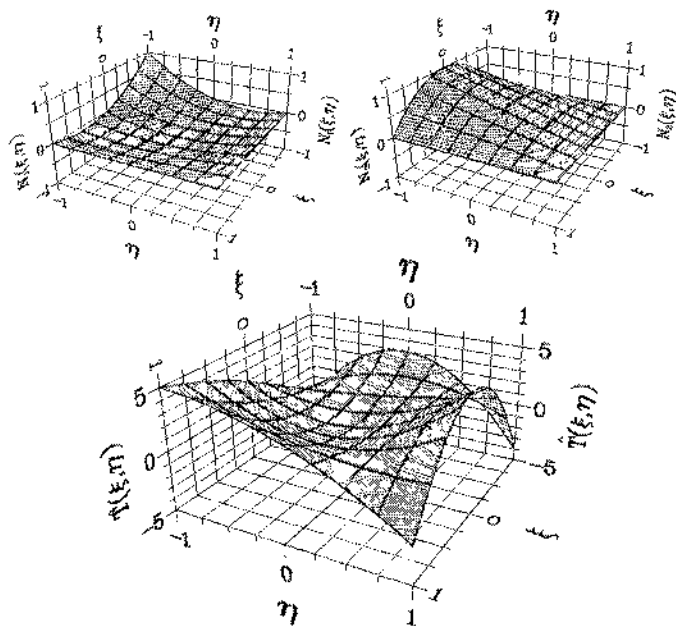


Figure 1.6.9 Form of shape functions and function for eight-noded quadrilateral.

The form of these shape functions is illustrated in Figure 1.6.9, which also shows the form of the function $\hat{T}$ when nodal values $T_1, T_2, \dots, T_8$ are prescribed.

A single eight-noded element can be used to obtain an approximate solution to the problem of equations (1.6.9)–(1.6.11). By selecting eight points (x_j, y_j) ; $j = 1, 2, \dots, 8$, on the boundary Γ , with (x_1, y_1) and (x_3, y_3) located at the end points of Γ_1 and (x_2, y_2) also located on Γ_1 , an element can be created which approximates the region Ω by mapping the region $-1 \leq \xi, \eta \leq 1$ by the isoparametric relation

$$x = \sum_{j=1}^8 x_j N_j(\xi, \eta) \quad y = \sum_{j=1}^8 y_j N_j(\xi, \eta) \quad (1.6.21)$$

When this element is used to construct a Galerkin approximate solution to problem (1.6.9)–(1.6.11), the trial and weighting function sets are

$$\mathcal{F}^{(5)} = \left\{ \hat{T}(x, y) | \hat{T}(x, y) = \psi + \sum_{j=4}^8 T_j N_j(\xi, \eta); \psi = \sum_{j=1}^3 f_j N_j(\xi, \eta); \hat{T} = \hat{f} \text{ on } \eta = -1 \right\} \quad (1.6.22)$$

$$\mathcal{W}^{(5)} = \left\{ W(x, y) | W(x, y) = \sum_{j=4}^8 b_j N_j(\xi, \eta); W = 0 \text{ on } \eta = -1 \right\} \quad (1.6.23)$$

with the relation between the coordinates (x, y) and (ξ, η) now being defined by the mapping of equation (1.6.21).

1.7 A PRACTICAL IMPLEMENTATION OF THE GALERKIN FINITE ELEMENT METHOD

In the description of the Galerkin finite element method which has been presented above, it can be observed that the nodes on Γ_1 , at which essential boundary conditions have to be imposed, are treated from the outset in a different fashion to the other nodes. This observation leads to unnecessary complexities when attempting to develop a computational implementation of the procedure.

To illustrate how these complexities can be overcome, consider the problem of solving the differential equation (1.6.1) using the element e shown in Figure 1.6.1. We begin by assuming that neither of the boundary conditions will be treated as essential conditions and work with the trial function set

$$\mathcal{T}^{(M+1)^2} = \left\{ \hat{T} \mid \hat{T} = \sum_{j=1}^{(M+1)^2} T_j N_j \right\} \quad (1.7.1)$$

and the weighting function set

$$\mathcal{W}^{(M+1)^2} = \left\{ W \mid W = \sum_{j=1}^{(M+1)^2} b_j N_j \right\} \quad (1.7.2)$$

An appropriate starting point for the construction of a finite element approximate solution is then to employ equation (1.3.12) as the basis of a variational formulation in the form: find $\hat{T}$ in $\mathcal{T}^{(M+1)^2}$ such that

$$\begin{aligned} \int_{\Omega} k \operatorname{grad} \hat{T} \cdot \operatorname{grad} N_i d\Omega &= \int_{\Omega} Q N_i d\Omega - \int_{\Gamma_1} N_i \hat{q} d\Gamma - \int_{\Gamma_2} N_i \aleph d\Gamma \\ &+ \Psi \int_{\Gamma_1} [\hat{T} - \hat{f}] N_i d\Gamma \end{aligned} \quad (1.7.3)$$

where $\hat{q}$ is an approximation to the unknown value of the flux on Γ_1 . Inserting the assumed form for $\hat{T}$ from equation (1.7.1), it follows that the finite element approximation is determined by the solution of the $(M+1) \times (M+1)$ matrix equation system

$$\mathbf{K}_e \mathbf{T}_e = \mathbf{r}_e - \mathbf{q}_e \quad (1.7.4)$$

where the vectors $\mathbf{T}_e$ and $\mathbf{q}_e$ are defined by

$$\mathbf{T}_e^T = (T_1, T_2, \dots, T_{(M+1)^2}) \quad \mathbf{q}_e^T = (q_1, q_2, \dots, q_{4M}, 0, \dots, 0) \quad (1.7.5)$$

and where

$$q_i = \int_{\Gamma} \hat{q} N_i d\Gamma \quad (1.7.6)$$

where $\hat{q}$ is now to be regarded as equal to the unknown heat flux on Γ_1 and equal to the specified heat flux (i.e. $\hat{q} = \aleph$) on Γ_2 . The matrix $\mathbf{K}_e$ and the vector $\mathbf{r}_e$ have

typical entries

$$\begin{aligned} K_{ij} &= \int_{\Omega} k \text{grad } N_j \cdot \text{grad } N_i \, d\Omega \\ r_i &= \int_{\Omega} Q N_i \, d\Omega + \Psi \int_{\Gamma_1} [\hat{T} - \hat{f}] N_i \, d\Gamma \end{aligned} \quad (1.7.7)$$

At this stage, we satisfy the boundary condition on Γ_1 by imposing the values

$$T_i = f_i \quad \text{for } i = 1, 2, \dots, M+1 \quad (1.7.8)$$

so that the integral over Γ_1 in equation (1.7.7) disappears from the formulation. Equation system (1.7.4) can then be solved for the approximations to the values of the temperature at the $M(M+1)$ nodes not subject to the specified temperature boundary condition. It is readily demonstrated that these will be exactly the values determined by the direct application of the Galerkin method as outlined previously. At the same time, the solution process will give values for the approximation to the integrated flux, q_i , at all nodes at which the value of the temperature has been specified.

It can be observed that, if Q is represented by its Lagrange interpolant over Ω and if N is represented by its Lagrange interpolant on Γ_2 , the approximations

$$\int_{\Omega} Q N_i \, d\Omega \approx \sum_{j=1}^{(M+1)^2} \left[\int_{\Omega} N_j N_i \, d\Omega \right] Q_j \quad \text{for } i = 1, 2, \dots, (M+1)^2 \quad (1.7.9)$$

and

$$\int_{\Gamma_2} N N_i \, d\Gamma \approx \sum_{j=M+1}^{4M} \left[\int_{\Gamma_2} N_j N_i \, d\Gamma \right] N_j \quad \text{for } i = M+1, M+2, \dots, 4M \quad (1.7.10)$$

can be employed to simplify the evaluation of the components of the vectors r defined in equation (1.7.7) and q in equation (1.7.6).

1.8 FURTHER OBSERVATIONS

We have now laid some of the important foundations related to the use of the finite element method for the approximate solution of heat conduction problems and we are ready to begin studying, in detail, the practical implementation of the method. However, before we embark on this study, a few further observations can be made on the theory which has been introduced in this chapter.

The reader will have observed that the shape functions for the two-dimensional linear and quadratic elements, introduced in the last section, contain additional terms to those that would be present in a complete linear or quadratic polynomial in ξ and η . Thus, for example, the linear element shape functions contain the

bilinear term $\xi\eta$, while the quadratic Lagrangian element shape functions contain the terms $\xi^2\eta$, $\xi\eta^2$ and $\xi^2\eta^2$. The quadratic serendipity element contains the additional terms $\xi^2\eta$ and $\xi\eta^2$. These parasitic polynomial terms cannot be removed from the nodal shape functions for square elements. The construction of elements which are triangular in shape generally leads to nodal shape functions which are in the form of complete polynomials, but such elements will not be treated in any detail in this book.

It will also be apparent that the evaluation of the entries in the matrix $\mathbf{K}_e$ and the vector $\mathbf{r}_e$ of equation (1.7.7) will require the evaluation of integrals, over the approximated area of interest, when the solution of a problem such as (1.6.9)–(1.6.11) is attempted by the Galerkin method. We will be following an approach in which such integrals are evaluated approximately, by standard numerical integration techniques in the mapped (ξ, η) plane. This topic will be addressed in the next chapter.

The analyst who is interested in improving the accuracy of the computation can achieve this by increasing the dimension M of the trial function set, i.e. by employing a higher-order element. In practice, however, the approach which is normally adopted is to use low-order elements (such as the linear or quadratic elements introduced above) and to improve the solution quality by sub-dividing the region of interest into a number of elements, as illustrated for example in Figure 1.2.3, and to employ a piecewise (local) low-order interpolation of the solution over each element. The resulting approximate solution $\hat{T}$ will possess C^0 continuity, e.g. considering an assembly of quadratic serendipity elements, inter-element continuity of the approximation is assured as the three nodes located along any element side determine a unique quadratic expansion. This sub-division of the region of interest means that the process of evaluating integrals, such as those appearing in equation (1.7.7), will require further attention. The way in which this complication can be handled will be considered in the next chapter.

ADDITIONAL READING

- Becker E. B., Carey G. F. and Oden J. T., *Finite Elements—An Introduction*, Volume 1, Prentice-Hall (1981).
- Carslaw H. S. and Jaeger J. C., *Conduction of Heat in Solids* (2nd Edition), Oxford University Press (1959).
- Johnson C., *Numerical Solution of Partial Differential Equations by the Finite Element Method*, Cambridge University Press (1987).
- Kreyszig E., *Advanced Engineering Mathematics*, 7th Edition, John Wiley & Sons (1993).
- Luikov A. V., *Analytical Heat Diffusion Theory*, Academic Press (1968).
- Milne R. D., *Applied Functional Analysis—An Introductory Treatment*, Pitman Publishing (1980).
- Prenter P. M., *Splines and Variational Methods*, Wiley (1975).

- Reddy J. N., *An Introduction to the Finite Element Method*, McGraw-Hill (1984).
- Strang G. and Fix G., *An Analysis of the Finite Element Method*, Prentice-Hall (1973).
- Zienkiewicz O. C. and Morgan K., *Finite Elements and Approximation*, John Wiley & Sons (1983).
- Zienkiewicz O. C. and Taylor R. L., *The Finite Element Method, 4th Edition*, Volume 1, McGraw-Hill (1989).