



Galerkin schemes for elliptic boundary value problems
by Bruce Victor Riley

A thesis submitted in partial fulfillment of the requirements for the degree of DOCTOR OF
PHILOSOPHY in Mathematics
Montana State University
© Copyright by Bruce Victor Riley (1982)

Abstract:

The Bubnov-Galerkin method is applied to two elliptic boundary value problems which arise in the mathematical model of two well known physical phenomena. In each instance, it is shown that the pertinent physical quantity being approximated can be feasibly obtained (accurately and economically) with the proposed procedures.

The sinc-collocation method is developed to numerically solve the radial Schrodinger eigenvalue problem. This Sturm-Liouville problem is classified as a singular problem for two distinct reasons.

Firstly, there is the infinite domain R^+ over which the solution is sought and, in instances, the eigenfunctions are singular at the origin. It is well known that these types of singularities pose difficulties for standard numerical procedures. It is shown that the sinc-collocation method surmounts these difficulties for the Schrodinger problem. Indeed, the method is applicable to Sturm-Liouville problems defined on infinite intervals in general. With regard to singular eigenfunctions, the error bound $O(N \exp\{-aN^{1/2}\})$, where $2N + 1$ basis functions are used, obtained for the eigenvalue approximations are valid whether or not the eigenfunctions are singular at the origin.

The annular strip method is developed to address a heat transfer problem defined on an annular domain Ω . The curved boundaries of Ω pose difficulties for the finite element method, and the accuracy of the resulting approximation will, in general, not be satisfactory relative to the work required for the acquisition of the approximation. The annular strip method is a viable alternative to the finite element method in that the former retains the salient features of the finite element method and avoids the boundary difficulties via a judicious choice of trial functions. The convergence results obtained for the annular strip method can be extended to applications of the finite strip method of the engineering literature.

GALERKIN SCHEMES FOR ELLIPTIC BOUNDARY VALUE PROBLEMS

by

BRUCE VICTOR RILEY

A thesis submitted in partial fulfillment
of the requirements for the degree

of

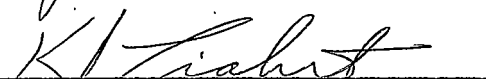
DOCTOR OF PHILOSOPHY

in

Mathematics

Approved:


Chairperson, Graduate Committee


Head, Major Department


Graduate Dean

MONTANA STATE UNIVERSITY
Bozeman, Montana

May, 1982

ACKNOWLEDGMENT

I wish to thank my advisor Dr. John R. Lund for his guidance and assistance throughout my association with him. In particular, his criticisms and suggestions during the writing of this thesis have been invaluable. My appreciation is also extended to Professors Norman Eggert, Donald Taylor and Russell Walker for their careful reading of this manuscript.

Thanks are due to NORCUS for their financial support during the summer of 1980, which enabled me to make great progress in my research, and to the MSU Research-Creativity Committee for funding part of the computer time required for the completion of my investigations.

Finally, I wish to thank my wife Peggie for her patience and encouragement during these past four years and for her skillful typing of this thesis.

TABLE OF CONTENTS

	<u>Page</u>
1. Introduction	1
1.1 Galerkin Schemes	1
1.2 One Dimensional Problem	5
1.3 Two Dimensional Problem	7
2. Preliminary Results	12
2.1 Introduction	12
2.2 Sobolev Spaces	13
2.3 Sinc Function Approximation	26
2.4 Perturbed Matrices	43
3. The Sinc-Collocation Method	47
3.1 Introduction	47
3.2 Finite Element Approximation	48
3.3 Sinc-Collocation Approximation	53
3.4 Convergence of the Eigenvalues	57
3.5 Numerical Implementation and Examples	64
4. The Annular Strip Method	78
4.1 Introduction	78
4.2 The Ritz-Galerkin Technique	83
4.3 Finite Element Approximation	91
4.4 Annular Strip Approximation	94

	<u>Page</u>
4.5 Numerical Implementation and Examples	113
5. Summary	125
Bibliography	131

LIST OF TABLES

<u>Table</u>	<u>Page</u>
3.1 Approximate eigenvalues for the radial Schrodinger equation with a harmonic oscillator potential	68
3.2 Approximate eigenvalues for the radial Schrodinger equation with a Wood-Saxon potential	72
3.3 Approximate eigenvalues for a radial Schrodinger equation with singular eigenfunctions	74
3.4 Approximate eigenvalues for the Hermite differential equation	76
4.1 Annular strip approximation to the solution of Example 4.9	118
4.2 Annular strip approximation to the solution of Example 4.10	123

LIST OF FIGURES

<u>Figure</u>	<u>Page</u>
1.1 The annulus $\Omega = \{(x, y) : R_1 < \sqrt{x^2 + y^2} < R_2\}$	9
2.1 The domains D and D_d	36
2.2 The domain $D = \{z : \arg z < d\}$	37
2.3 The domain $D = \{z : \arg \sinh z < d\}$	38
4.1 Fuel rod partially inserted into core channel	79
4.2 Cross section of the fuel rod within the core channel	79
4.3 Fuel rod partially inserted into a bowed core channel	81
4.4 Cross section of the fuel rod within the bowed core channel	81
4.5 The domain Ω with inner boundary $\partial_1 \Omega$ and outer boundary $\partial_2 \Omega$	82
4.6 The approximation of the annulus Ω by the polygonal domain Ω_h	92
4.7 A partial triangularization of the domain bounded by $\partial_2 \Omega_h$	92
4.8 Ω partitioned into annular strips	97
4.9 The Chapeau function Λ_i	98
4.10 An example of the algebraic system for the determination of an annular strip approximation	116

ABSTRACT

The Bubnov-Galerkin method is applied to two elliptic boundary value problems which arise in the mathematical model of two well known physical phenomena. In each instance, it is shown that the pertinent physical quantity being approximated can be feasibly obtained (accurately and economically) with the proposed procedures.

The sinc-collocation method is developed to numerically solve the radial Schrodinger eigenvalue problem. This Sturm-Liouville problem is classified as a singular problem for two distinct reasons. Firstly, there is the infinite domain $\mathbb{R}_+$ over which the solution is sought and, in instances, the eigenfunctions are singular at the origin. It is well known that these types of singularities pose difficulties for standard numerical procedures. It is shown that the sinc-collocation method surmounts these difficulties for the Schrodinger problem. Indeed, the method is applicable to Sturm-Liouville problems defined on infinite intervals in general. With regard to singular eigenfunctions, the error bound $O(N \exp\{-\alpha N^{1/2}\})$, where $2N + 1$ basis functions are used, obtained for the eigenvalue approximations are valid whether or not the eigenfunctions are singular at the origin.

The annular strip method is developed to address a heat transfer problem defined on an annular domain Ω . The curved boundaries of Ω pose difficulties for the finite element method, and the accuracy of the resulting approximation will, in general, not be satisfactory relative to the work required for the acquisition of the approximation. The annular strip method is a viable alternative to the finite element method in that the former retains the salient features of the finite element method and avoids the boundary difficulties via a judicious choice of trial functions. The convergence results obtained for the annular strip method can be extended to applications of the finite strip method of the engineering literature.

1. INTRODUCTION

1.1 Galerkin Schemes

Many boundary value problems that describe various physical phenomena, indeed, many physical problems in general, have solutions that satisfy a variational principle. Of the many numerical techniques used to find solutions of such problems, the finite element method has established itself as a viable and competitive alternative to the classical finite difference methods for boundary value problems. The origins of the finite element method may be traced back to the work of Walter Ritz [18] in his pioneering work on various structural analysis problems. The method is still widely used for these problems today, and moreover, the finite element method has been successfully applied to problems in heat conduction, fluid dynamics and potential (electric and magnetic) field theory.

In this thesis, the Bubnov-Galerkin method, a general procedure for which the finite element method is a specific case, is applied to a class of elliptic differential equations. The Galerkin scheme may be described as follows: Assume Λ is an elliptic differential operator and it is desired to find a solution u to the problem

$$(1.1) \quad \Lambda u = f \quad \text{on } \Omega \subset \mathbb{R}^n,$$

where (1.1) is equipped with appropriate boundary conditions. It is assumed that $\Lambda : H \longrightarrow H$ where H is an appropriate Hilbert space. For each positive integer n , let $\{\phi_i\}_{i=1}^n$ be an independent set in H and define a trial solution to (1.1) by the formula

$$(1.2) \quad u_n = \sum_{i=1}^n \alpha_i \phi_i.$$

The α_i are determined by the requirement of orthogonality of the error:

$$\Lambda u_n - f \perp \phi_j \quad \text{for } j = 1, \dots, n.$$

Thus, the α_i may be computed from the discrete problem

$$(1.3) \quad \sum_{i=1}^n \alpha_i (\Lambda \phi_i, \phi_j) = (f, \phi_j) \quad \text{for } j = 1, \dots, n.$$

Different Galerkin approximations are obtained according to the selection of the basis $\{\phi_i\}$ which is made. This selection is influenced by the problem being solved, the domain on which the problem is defined and, to a degree, on personal preference or prejudice.

In the finite element method the trial functions ϕ_i are usually piecewise polynomials. A given domain Ω is represented as a finite

collection of geometrically simpler subdomains Ω_k (finite elements) in such a manner that

$$\Omega \approx \bigcup_{k=1}^m \Omega_k.$$

The Ω_k may share common boundaries but are otherwise disjoint. A trial function ϕ_i is chosen so that its support is contained in only a few of the Ω_k , and over these subdomains ϕ_i is pieced together from polynomials of low degree. With the small support, relative to all of the domain Ω , for each ϕ_i , many of the terms in (1.3) are zero, and consequently, the system in (1.3) to be solved is a sparse system of equations. It is this choice of trial functions, and the resulting sparse system of equations (1.3) to be solved, which is largely responsible for the success of the finite element method.

Although the finite element method is highly adaptable there are many problems to which the method does not directly apply. This thesis addresses two of these problems. The first of these problems is motivated by a study of the classical radial Schrodinger equation and the methods developed apply directly to the case of a singular Sturm-Liouville problem, in particular, a Sturm-Liouville problem that is defined on an infinite or semi-infinite interval. The Schrodinger problem arises from an elliptic partial differential

equation in $\mathbb{R}^3$. In the important axil-symmetric case, the partial differential operator admits a separation of variables in spherical coordinates of which the two angular solutions are well known and the remaining radial portion is the singular Sturm-Liouville problem. The second problem arises again as a domain problem, but this time for the Poisson equation in $\mathbb{R}^2$. Here the trouble is not one of an infinite domain, indeed the domain is compact, but due to its symmetries the classical selection of subdomains Ω_k (rectangles or triangles) boundary errors are introduced. In each problem the given domain Ω must first be replaced by an approximate domain $\bigcup_{k=1}^m \Omega_k$. The finite element method is then applied to the problem on the approximate domain $\bigcup_{k=1}^m \Omega_k$. In general, it is expected that as m increases the approximate domain approaches Ω and more accurate approximations to the true solution are obtained. However, requiring $\bigcup_{k=1}^m \Omega_k$ to be close to Ω does not guarantee an accurate approximation to the true solution of the problem. There is an example discussed in Strang and Fix [29] which illustrates that the approximate solutions may actually fail to converge to the solution of the original problem. Even in the case that $\bigcup_{k=1}^m \Omega_k$ approaches Ω and the solution of the discrete problem approaches that of $\Delta u = f$, the increasing problem size and resulting computational costs may be prohibitively large.

1.2 One Dimensional Problem

Atomic and nuclear physics investigations frequently require the calculation of the eigenvalues of the radial Schrodinger equation

$$(1.4) \quad -u''(x) + q(x)u(x) = \lambda p(x)u(x) \quad x \in \mathbb{R}_+ \quad (0, \infty)$$

subject to the boundary conditions

$$(1.5) \quad u(0) = 0 \quad \text{and} \quad \lim_{x \rightarrow \infty} u(x) = 0.$$

In most physical problems, the functions q and p are of the general form

$$q(x) = r(x) + s(x)/x^2$$

and

$$p(x) = t(x) + w(x)/x^2,$$

where r , s , t and w are bounded on every finite interval $[0, b]$. In Chapter 3 the singular Sturm-Liouville problem (1.4) and (1.5) is investigated.

The differential equation (1.4) with appropriate boundary conditions on a finite interval $[a, b]$ has been amply discussed and many satisfactory numerical schemes, mostly based on polynomial basis functions are available for the approximate solution of the problem. The situation for the aforementioned singular problem is in a much

less refined state. This is the case whether the singularity in the problem is due to the infinite domain ($\Omega = \mathbb{R}_+$) over which (1.4) is to be solved or whether the solution u of (1.4) is singular at a point in Ω .

The solution of Equation (1.4) via the finite element method, with bounded q and p on a finite interval, and subject to homogeneous boundary conditions, has been studied in a number of papers [5,11,12,32]. These investigations use polynomial spline function expansions of the eigenfunctions in the Rayleigh-Ritz method to calculate approximate eigenvalues. This technique, as well as its extension by Shore [22-24], seeks a solution of Equation (1.4) on a finite interval $[0,b]$ subject to the homogeneous boundary conditions $u(0) = u(b) = 0$ for large b . This method exchanges one sort of difficulty for another, that is, the singular nature of (1.4) due to the infinite domain has been traded for a truncation error due to the altered domain. It should, however, be pointed out that Schoombie and Botha [20] have recently obtained theoretical error estimates for the truncation error introduced by the altered domain.

An alternate numerical procedure is developed in Chapter 3 for the problem (1.4) and (1.5). The method is a Galerkin scheme as described by (1.2) and (1.3) where the basis functions ϕ_i are compositions of certain conformal maps with the sinc function. As the

latter are defined on all of $\mathbb{R}$, the truncation of the infinite interval, as discussed in the previous paragraph, is not required, and therefore, the error introduced by solving an approximate domain problem is avoided.

Another advantage of the method developed in Chapter 3 is due to the highly accurate approximation of the inner products in (1.3) which are obtainable with the sinc function basis elements. These quadrature rule approximations along with the pertinent sinc function differentiation formulas are contained in Chapter 2.

Chapter 2 also contains some results obtained on the perturbation of symmetric matrices. These results are needed in Chapter 3 to show the convergence of the eigenvalues of the discrete problem (1.3) to the eigenvalues of the continuous problem (1.4).

1.3 Two Dimensional Problem

In Chapter 4, the problem

$$(1.6) \quad \Delta T = \frac{\partial^2 T}{\partial x^2} + \frac{\partial^2 T}{\partial y^2} = f \quad \text{on } \Omega \subset \mathbb{R}^2$$

with boundary conditions for T on $\partial\Omega$ is considered. The techniques of Chapter 3 are directly applicable to the problem (1.6) for appropriate domains and conformal maps. For the case of the square or rectangle, see the examples in [26] and [27]. Indeed, if f in (1.6) is replaced

by $-\lambda T$, then by a separation of variables one is led to an eigenvalue problem analogous to the problem studied in Chapter 3 except that in the present case the domain is finite.

The investigations of Chapter 4 are motivated by a heat transfer problem which arises because of changes in an aging nuclear reactor. The domain Ω is as pictured in Figure 1.1 and the problem is to solve (1.6) with $f = 0$ on Ω subject to the boundary conditions

$$T = 0 \quad \text{on } \partial_1\Omega$$

and

$$\frac{dT}{dn} + gT = k_1g + k_2 \quad \text{on } \partial_2\Omega$$

where $\frac{dT}{dn}$ denotes the derivative of T in the direction of the outward normal to the curve $\partial_2\Omega$, g is a continuous nonnegative function defined on $\partial_2\Omega$, and k_i ($i = 1, 2$) are constants.

The equivalent variational form of this problem and the classical selection of trial functions for the finite element method are discussed in Chapter 4. The annulus Ω is first approximated by a triangularized polygonal domain $\bigcup_{k=1}^m \Omega_k$ and the trial functions ϕ_i are defined on this partitioned domain. The polygonal approximation of the circular boundaries in the finite element method produce two types of boundary errors. This is due to the different nature of the boundary conditions on $\partial_1\Omega$ and $\partial_2\Omega$. Both errors will be discussed in

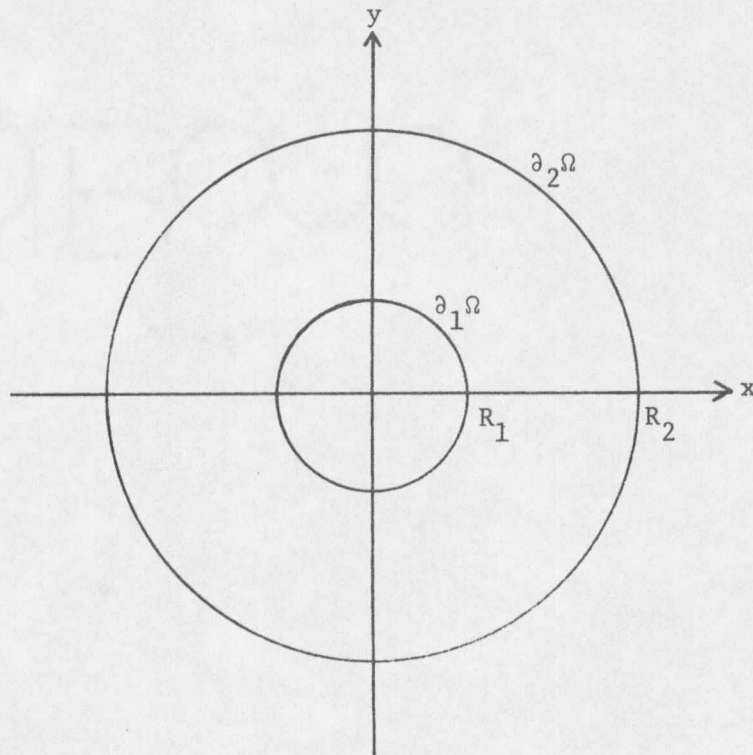


Figure 1.1 The annulus $\Omega = \{(x,y) = R_1 < \sqrt{x^2 + y^2} < R_2\}$. $\partial_1 \Omega$ and $\partial_2 \Omega$ denotes the inside and outside boundaries of Ω , respectively.

Chapter 4.

The numerical scheme developed in Chapter 4 is a semianalytic method. It, like the finite element method, is a variation of the Ritz-Galerkin technique. In order to avoid the error induced by approximating the boundaries of the annulus, Ω is partitioned into annular strips. The trial functions, unlike the finite element trial

functions which are piecewise polynomials in all directions, are products of trigonometric functions defined around the strips and piecewise polynomials defined across the strips. This choice of trial functions avoids the earlier boundary errors, yet retains a sparse discrete system (1.3). It is shown in Chapter 4 how to order the unknown coefficients of the basis function in (1.3) so that the resulting sparse system can be efficiently solved. The matrix of the system (1.3) is $(mM) \times (mM)$ where M is the number of trigonometric functions used and m is the number of subintervals into which the radial interval is partitioned. Instead of the usual $O((mM)^3)$ operations in the general Gaussian elimination procedure to find the unknown coefficients, matters are arranged so that the operation count is on the order of $O((m-1)M + M^3)$. This, of course, significantly reduces the computer time and the examples indicate accuracy comparable to that obtained by the finite element method.

The scheme developed in Chapter 4 is an example of what the engineers call the finite strip method [10]. The strip method was developed by structural analysts to reduce the high cost of analyzing regular geometric structures. For many years the strip method appeared only in the engineering literature connected with structural analysis. Not until recently, when Chakrabarti [9] used the strip method for his study of heat conduction in rectangular plates, has the

method appeared in a field of study outside that of structural analysis. The ideas of the strip method appear to be applicable to many problems, and therefore, a thorough analysis of the method is warranted.

In this direction, convergence results for the aforementioned annular strip method are obtained. There have appeared many examples in the literature [9,10,30] which give numerical evidence of the convergence of the finite strip method. However, this author is not aware of any other results showing that the strip approximates do indeed converge to the true solution. This author believes that the results of Chapter 4 can be extended to other applications (thermal stress and structural analysis) of the finite strip method, and therefore, the method merits investigation in order to place it on a firm mathematical foundation.

2. PRELIMINARY RESULTS

2.1 Introduction

This chapter contains a number of results from Hilbert space theory, sinc function approximation, and linear algebra which are required in later chapters of the thesis. Most of the results are stated without proof but are equipped with appropriate references.

The finite element method, as it applies to elliptic boundary value problems, is based on Hilbert space theory, and in particular, the notions of Sobolev spaces, traces, and interpolation spaces. A brief discussion of Sobolev spaces is given in Section 2.2. The material can be found in the works of Adams [1], Agmon [2] and Aubin [3].

In Section 2.3, various results on sinc function approximation are given. These results have been extracted from the works of Stenger, his colleagues and students [14,15,25-27]. The survey paper [27] is an excellent source for the numerical methods based on the sinc function.

Section 2.4 contains various theorems on the symmetric eigenvalue problem which can be found in Steward [19]. These results are used to obtain estimates for the spectral radius of perturbed symmetric matrices.

2.2 Sobolev Spaces

Sobolev spaces are Banach spaces of weakly differentiable functions of several real variables. They arise in connection with numerous problems in the theory of partial differential equations and, in particular, elliptic equations. The Sobolev spaces also provide a natural setting for the mathematical foundation of the finite element method. There are at least two ways of constructing the Sobolev spaces and since a knowledge of both constructions lends useful insight to an understanding of the variational principles in Chapters 3 and 4, both developments are herein reviewed. The following notation will be used in both developments.

Let Ω be an open domain in n -dimensional euclidean space $\mathbb{R}^n$, the points of which are denoted by $x = (x_1, x_2, \dots, x_n)$, and let $\partial\Omega$ denote the boundary of Ω . For the n -tuple $\alpha = (\alpha_1, \alpha_2, \dots, \alpha_n)$ of nonnegative integers define

$$|\alpha| = \alpha_1 + \alpha_2 + \dots + \alpha_n,$$

and for $x \in \mathbb{R}^n$ and $\alpha = (\alpha_1, \alpha_2, \dots, \alpha_n)$ define

$$x^\alpha = x_1^{\alpha_1} x_2^{\alpha_2} \dots x_n^{\alpha_n}$$

In particular, this notation is convenient for the compact expression of the differential operator

$$D^{\alpha} = D_1^{\alpha_1} D_2^{\alpha_2} \dots D_n^{\alpha_n} = \frac{\partial^{|\alpha|}}{\partial x_1^{\alpha_1} \partial x_2^{\alpha_2} \dots \partial x_n^{\alpha_n}}$$

Consider a linear differential operator of order m

$$A(x, D) = \sum_{|\alpha| \leq m} a_{\alpha}(x) D^{\alpha}$$

where the a_{α} are real valued functions defined on some open domain Ω in $\mathbb{R}^n$. It is assumed that not all coefficients a_{α} with $|\alpha| = m$ are zero on Ω . With each operator $A(x, D)$ there is associated the homogeneous polynomial

$$A'(x, \xi) = \sum_{|\alpha|=m} a_{\alpha}(x) \xi^{\alpha}$$

of degree m in the variable $\xi \in \mathbb{R}^n$. The corresponding operator $A'(x, D)$ is called the principle part of $A(x, D)$. The operator $A(x, D)$ is said to be elliptic at a point x_0 if and only if for any $\xi \neq 0$, $A'(x_0, \xi) \neq 0$.

For example, in Chapter 3 the Sturm-Liouville operator

$$A_1(x, D) = -D^2 + q(x)$$

is investigated for functions defined on $\mathbb{R}_+ = (0, \infty)$. The principle part of $A_1(x, D)$ is

$$A_1'(x, D) = -D^2.$$

It is clear that $A_1(x, D)$ is an elliptic operator. In Chapter 4 the Laplacian operator

$$A_2(x, D) = D_1^2 + D_2^2$$

is considered for functions defined on an annulus in $\mathbb{R}^2$. The principle part of $A_2(x, D)$ is itself and it is again clear that $A_2(x, D)$ is an elliptic operator.

Now, for an open bounded domain $\Omega \subset \mathbb{R}^n$, let $L^2(\Omega)$ denote the usual space of square integrable real valued functions over Ω endowed with the inner product

$$(u, v) = \int_{\Omega} uv \, dx.$$

For any nonnegative integer m , let $C^m(\Omega)$ be the class of m times continuously differentiable functions on Ω . Define

$$C^\infty(\Omega) = \bigcap_{m=1}^{\infty} C^m(\Omega)$$

and let $C_0^\infty(\Omega)$ denote the subset of $C^\infty(\Omega)$ consisting of functions having compact support contained in Ω . $C_0^\infty(\Omega)$ is called the set of

test functions for Ω .

The first construction of a Sobolev space follows the development given by Agmon [2]. For each $u \in C^m(\Omega)$ define

$$(2.1) \quad |||u|||_m = \left(\int_{\Omega} \sum_{|\alpha| \leq m} (D^\alpha u)^2 dx \right)^{1/2}$$

and let $C_f^m(\Omega)$ denote the subset of $C^m(\Omega)$ consisting of the functions u for which $|||u|||_m < \infty$. For $u, v \in C_f^m(\Omega)$, let

$$(2.2) \quad (u, v)_m = \int_{\Omega} \sum_{|\alpha| \leq m} (D^\alpha u) (D^\alpha v) dx.$$

For example, if $\Omega \subset \mathbb{R}^2$ and $u, v \in C_f^1(\Omega)$, then

$$(2.3) \quad |||u|||_1 = \left(\iint_{\Omega} (u^2 + (\frac{\partial u}{\partial x})^2 + (\frac{\partial u}{\partial y})^2) dx dy \right)^{1/2}.$$

and

$$(2.4) \quad (u, v)_1 = \iint_{\Omega} (uv + \frac{\partial u}{\partial x} \frac{\partial v}{\partial x} + \frac{\partial u}{\partial y} \frac{\partial v}{\partial y}) dx dy$$

Note that $C_f^m(\Omega)$ is an inner product space with the inner product $(\cdot, \cdot)_m$. Let $H^m(\Omega)$ be the completion of $C_f^m(\Omega)$ with respect to $|||\cdot|||_m$. The Hilbert space $H^m(\Omega)$ is called a Sobolev space.

The elements of $H^m(\Omega)$ may be thought of as follows: Let $\{u_k\}$ be a Cauchy sequence in $C_F^m(\Omega)$ with respect to the Sobolev norm (2.1). Then for each α , $|\alpha| \leq m$ the inequality

$$\left(\int_{\Omega} (D^{\alpha} u_k - D^{\alpha} u_{\ell})^2 dx \right)^{1/2} \leq \|u_k - u_{\ell}\|_m$$

implies that $\{D^{\alpha} u_k\}$ is a Cauchy sequence in $L^2(\Omega)$. Now, since $L^2(\Omega)$ is complete, there is a $u^{\alpha} \in L^2(\Omega)$ such that $\{D^{\alpha} u_k\}$ converges to u^{α} in $L^2(\Omega)$. Therefore, an element $u \in H^m(\Omega)$ is an L^2 -function over Ω with the property that there exists a sequence $\{u_k\} \subset C_F^m(\Omega)$ such that $\{D^{\alpha} u_k\}$ is a Cauchy sequence in $L^2(\Omega)$ for all α , $|\alpha| \leq m$, and u_k converges to u in $L^2(\Omega)$.

The limit function u^{α} is called the strong L^2 -derivative of order α of the function $u \in H^m(\Omega)$, and is unique up to sets of measure zero. Hence, if $u \in C_F^m(\Omega)$, then $u \in H^m(\Omega)$ and the strong L^2 -derivatives u^{α} of u are exactly $D^{\alpha} u$. This is easily seen by choosing the Cauchy sequence $\{u_k\}$ associated with u as $u_k = u$ for each k .

The second construction of a Sobolev space follows the development given by Adams [1]. Consider the set of test functions $C_0^{\infty}(\Omega)$. A sequence $\{\phi_k\} \subset C_0^{\infty}(\Omega)$ converges to $\phi_0 \in C_0^{\infty}(\Omega)$ if and only if there is a compact set $K \subset \Omega$ such that the support of each ϕ_k ($k = 0, 1, \dots$) is contained in K and $\{D^{\alpha} \phi_k\}$ converges uniformly to $D^{\alpha} \phi_0$ uniformly for all α .

A linear functional F on $C_0^\infty(\Omega)$ is called a distribution if it is continuous in the sense that if $\{\phi_k\}$ converges to ϕ_0 in $C_0^\infty(\Omega)$, then $\{F(\phi_k)\}$ converges to $F(\phi_0)$ in $\mathbb{R}$. An example of a distribution is the linear functional F_u induced by a function $u \in L^2(\Omega)$, where

$$F_u(\phi) = \int_{\Omega} u\phi \, dx \quad \text{for all } \phi \in C_0^\infty(\Omega).$$

A second example is the dirac mass δ_p concentrated at the point $p \in \Omega$ defined by

$$\delta_p(\phi) = \phi(p) \quad \text{for all } \phi \in C_0^\infty(\Omega).$$

Let $C_0^\infty(\Omega)^*$ denote the set of distributions on $C_0^\infty(\Omega)$. $C_0^\infty(\Omega)^*$ is a vector space under the usual definitions of addition

$$(F + G)(\phi) = F(\phi) + G(\phi)$$

and scalar multiplication

$$(aF)(\phi) = a(F(\phi)),$$

where $F, G \in C_0^\infty(\Omega)^*$, $\phi \in C_0^\infty(\Omega)$, and $a \in \mathbb{R}$. $C_0^\infty(\Omega)^*$ is topologized with the weak-* topology. That is, the sequence $\{F_k\}$ converges to F_0 in $C_0^\infty(\Omega)^*$ if and only if $F_k(\phi)$ converges to $F_0(\phi)$ in $\mathbb{R}$ for all $\phi \in C_0^\infty(\Omega)$.

The following observation motivates the definition of the derivative of a distribution. Observe that for $u \in C^m(\Omega)$ and

$\phi \in C_0^\infty(\Omega)$, $|\alpha|$ integration by parts yield the identity

$$(2.5) \quad \int_{\Omega} (D^\alpha u) \phi \, dx = (-1)^{|\alpha|} \int_{\Omega} u (D^\alpha \phi) \, dx$$

for each α , $|\alpha| \leq m$. Equation (2.5) is used to define the derivatives of the distribution F_u induced by u . The identity

$$D^\alpha F_u = F_{D^\alpha u}$$

follows from Equation (2.5) and the sequence of equalities

$$\begin{aligned} F_{D^\alpha u}(\phi) &= \int_{\Omega} (D^\alpha u) \phi \, dx \\ &= (-1)^{|\alpha|} \int_{\Omega} u (D^\alpha \phi) \, dx \\ &= (-1)^{|\alpha|} F_u(D^\alpha \phi) \quad \text{for all } \phi \in C_0^\infty(\Omega). \end{aligned}$$

Hence, the derivative of order α for any distribution F is defined by

$$D^\alpha F(\phi) = (-1)^{|\alpha|} F(D^\alpha \phi) \quad \text{for all } \phi \in C_0^\infty(\Omega).$$

With this definition, $D^\alpha F$ is indeed a distribution, and in fact, each D^α is a continuous linear operator on $C_0^\infty(\Omega)^*$. Therefore, every distribution in $C_0^\infty(\Omega)^*$ has derivatives of all orders in $C_0^\infty(\Omega)^*$.

An important class of distributions are those distributions which are induced by functions in $L^2(\Omega)$. Let $u \in L^2(\Omega)$, if there is a function $v \in L^2(\Omega)$ such that

$$D^{\alpha}F_u(\phi) = F_v(\phi) \quad \text{for all } \phi \in C_0^{\infty}(\Omega),$$

then it is unique up to sets of measure zero and it is called the weak L^2 -derivative of u . It is apparent that if u has a continuous derivative $D^{\alpha}u$ in the classical sense, then $D^{\alpha}u$ is also a weak L^2 -derivative of u . Thus, the notation $D^{\alpha}u$ will be used to denote the weak L^2 -derivative of u or the classical derivative of u depending on the context.

Now, for each nonnegative integer m , let $W^m(\Omega)$ denote the class of functions from $L^2(\Omega)$ which have weak L^2 -derivatives of all orders α , $|\alpha| \leq m$. Note that the Sobolev inner product (2.2) is an inner product on $W^m(\Omega)$. In fact, $W^m(\Omega)$ is complete with respect to the Sobolev norm. Therefore, $W^m(\Omega)$ is also called a Sobolev space and a result due to Meyers and Serrin [17] shows that $H^m(\Omega) = W^m(\Omega)$.

To summarize, there are two ways to consider a function $u \in H^m(\Omega)$. From the first construction, u is an L^2 -function with the property that there exists a sequence $\{u_k\} \subset C_0^m(\Omega)$ which is a Cauchy sequence with respect to the Sobolev norm (2.1), and u_k converges to u in $L^2(\Omega)$. From the second construction, u is an L^2 -function whose distributional derivatives $D^{\alpha}F_u$, $|\alpha| \leq m$, are equal to distributions F_v induced by L^2 -functions v .

The problems discussed in Chapters 3 and 4 are naturally posed in

the Sobolev space $H^1(\Omega)$ and the remainder of this section is devoted to a closer examination of $H^1(\Omega)$. The problems in these chapters are boundary value problems, and therefore, it is necessary to assign boundary values to the elements of $H^1(\Omega)$. The following development is taken from Aubin [3].

Let $\bar{\Omega}$ denote the closure of Ω in $\mathbb{R}^n$. The space $C^1(\bar{\Omega})$ is dense in $C^1(\Omega)$ with respect to the Sobolev norm (2.3), and therefore, dense in $H^1(\Omega)$ (recall the first construction of $H^1(\Omega)$). If $u \in C^1(\bar{\Omega})$, define the trace $\gamma_0 u$ of u to be the restriction of u to the boundary $\partial\Omega$ of Ω . The mapping

$$\gamma_0: C^1(\bar{\Omega}) \longrightarrow L^2(\partial\Omega)$$

is a continuous linear operator, where $C^1(\bar{\Omega})$ is given the Sobolev norm (2.3). The operator γ_0 can be extended to a continuous linear operator from $H^1(\Omega)$ into $L^2(\partial\Omega)$. Hence, boundary values have been assigned to each function $u \in H^1(\Omega)$. Denote by $H_0^1(\Omega)$ the class of functions $u \in H^1(\Omega)$ such that $\gamma_0 u = 0$. $H_0^1(\Omega)$ is a Sobolev space. It is the completion of $C_0^\infty(\Omega)$ with respect to the Sobolev norm (2.3), and is dense in $L^2(\Omega)$. Aubin shows that if Ω has a smooth boundary, then the trace operator γ_0 maps $H^1(\Omega)$ onto a dense subset of $L^2(\partial\Omega)$, and that the image S of $H^1(\Omega)$ under γ_0 is a Hilbert space when it is supplied with the strongest norm for which γ_0 is still continuous.

Finally, a second trace operator γ_1 from a subset of $H^1(\Omega)$ into S will be the generalized normal derivative of the function u from the appropriate subset of $H^1(\Omega)$. The following general Hilbert space development can also be found in [3].

Let U , V and W be Hilbert spaces, γ a continuous linear operator from V onto U , and $a(\cdot, \cdot)$ a continuous bilinear form on $V \times V$. Assume that V is contained in W with a stronger topology and that the kernel V_0 of γ is dense in W .

$$\begin{array}{ccccc} V_0 & \hookrightarrow & V & \hookrightarrow & W \\ & & \downarrow \gamma & & \\ & & U & & \end{array}$$

For any fixed $u \in V$, the linear map that associates the value $a(u, v)$ with $v \in V_0$ is continuous on V_0 . Denote by Au this linear functional on V_0 . That is, if V_0^* is the dual of V_0 , then $Au \in V_0^*$ is the functional defined by :

$$(Au, v)_V = a(u, v) \quad \text{for all } v \in V_0.$$

The operator A associating $Au \in V_0^*$ with any $u \in V$ is apparently a continuous linear operator from V into V_0^* . If $W = W^*$ in the sense of the Riesz theorem (the canonical isomorphism), then W is a dense

subspace of V_0^* . Define the set

$$W(A) = \{u \in V : Au \in W\}$$

and in the case that $W(A)$ is nonempty, the following generalized Green's formula is obtained.

Theorem 2.1. Suppose $U, V, W, W(A), \gamma$ and $a(\cdot, \cdot)$ satisfy the above hypothesis, then there is a unique operator γ_1 mapping $W(A)$ into V^* such that

$$(2.6) \quad a(u, v) = (Au, v) + \langle \gamma_1 u, \gamma v \rangle \quad \text{for any } u \in W(A) \text{ and } v \in V,$$

where $(\cdot, \cdot)$ is the inner product on W and $\langle \cdot, \cdot \rangle$ is the duality pairing on $U^* \times U$.

To obtain the generalized normal derivative alluded to earlier, consider the Laplacian operator

$$\Delta = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2}$$

for functions defined on an open bounded set Ω of $\mathbb{R}^2$. the classical Green's formula is

Theorem 2.2. Let u and v be twice continuously differentiable on the closed set $\bar{\Omega}$, then

$$(2.7) \quad \iint_{\Omega} \left(\frac{\partial u}{\partial x} \frac{\partial v}{\partial x} + \frac{\partial u}{\partial y} \frac{\partial v}{\partial y} \right) dx dy = \iint_{\Omega} (-\Delta u) v \, dx dy + \int_{\partial \Omega} \frac{du}{dn} v \, d\sigma.$$

Theorem 2.1, in particular, Equation (2.6), is the extension of Equation (2.7) to a class of functions in the Sobolev space $H^1(\Omega)$.

Identify the spaces in Theorem 2.1 as $V = H^1(\Omega)$ and $W = L^2(\Omega)$, and the operator γ as the trace operator

$$\gamma_0: H^1(\Omega) \longrightarrow L^2(\partial \Omega).$$

The kernel of γ_0 is $V_0 = H_0^1(\Omega)$ and $U = S$ is the image of γ_0 supplied with the strongest norm for which γ_0 is continuous.

Now define a bilinear form on $H^1(\Omega) \times H^1(\Omega)$ by

$$a(u, v) = \iint_{\Omega} \left(\frac{\partial u}{\partial x} \frac{\partial v}{\partial x} + \frac{\partial u}{\partial y} \frac{\partial v}{\partial y} \right) dx dy.$$

$a(u, v)$ is clearly bilinear and its continuity follows from an application of the Cauchy-Schwartz inequality in the calculation:

$$\begin{aligned} a(u, v) &= \iint_{\Omega} \left(\frac{\partial u}{\partial x} \frac{\partial v}{\partial x} + \frac{\partial u}{\partial y} \frac{\partial v}{\partial y} \right) dx dy \\ &\leq \left(\iint_{\Omega} \left(\frac{\partial u}{\partial x} \right)^2 dx dy \right)^{1/2} \left(\iint_{\Omega} \left(\frac{\partial v}{\partial x} \right)^2 dx dy \right)^{1/2} \\ &\quad + \left(\iint_{\Omega} \left(\frac{\partial u}{\partial y} \right)^2 dx dy \right)^{1/2} \left(\iint_{\Omega} \left(\frac{\partial v}{\partial y} \right)^2 dx dy \right)^{1/2} \end{aligned}$$

$$\leq 2 \|u\|_1 \|v\|_1.$$

Suppose $u \in H^1(\Omega)$ such that u has the weak L^2 -derivatives $\frac{\partial^2 u}{\partial x^2}$ and $\frac{\partial^2 u}{\partial y^2}$ (recall the second construction of a Sobolev space). Notice that $\frac{\partial^2 u}{\partial x^2}$ and $\frac{\partial^2 u}{\partial y^2}$ are weak L^2 -derivatives of $\frac{\partial u}{\partial x}$ and $\frac{\partial u}{\partial y}$, respectively. That is,

$$\iint_{\Omega} \frac{\partial^2 u}{\partial x^2} \phi \, dx dy = - \iint_{\Omega} \frac{\partial u}{\partial x} \frac{\partial \phi}{\partial x} \, dx dy$$

and

$$\iint_{\Omega} \frac{\partial^2 u}{\partial y^2} \phi \, dx dy = - \iint_{\Omega} \frac{\partial u}{\partial y} \frac{\partial \phi}{\partial y} \, dx dy$$

for all $\phi \in C_0^\infty(\Omega)$. Combining these two equations gives the equality

$$(2.8) \quad \iint_{\Omega} (-\Delta u) \phi \, dx dy = \iint_{\Omega} \left(\frac{\partial u}{\partial x} \frac{\partial \phi}{\partial x} + \frac{\partial u}{\partial y} \frac{\partial \phi}{\partial y} \right) \, dx dy$$

for all $\phi \in C_0^\infty(\Omega)$. Moreover, the density of $C_0^\infty(\Omega)$ in $H_0^1(\Omega)$ and the continuity of $a(\cdot, \cdot)$ shows that (2.8) persist for all $\phi \in H_0^1(\Omega)$.

It is apparent that $A = -\Delta$ is the formal operator associated with the bilinear form $a(\cdot, \cdot)$. The set

$$L^2(-\Delta) = \{u \in H^1(\Omega) : -\Delta u \in L^2(\Omega)\}$$

is clearly nonempty since it contains $C_0^\infty(\Omega)$.

The hypotheses of Theorem 2.1 are all satisfied and Equation (2.6) reads

$$\iint_{\Omega} \left(\frac{\partial u}{\partial x} \frac{\partial v}{\partial x} + \frac{\partial u}{\partial y} \frac{\partial v}{\partial y} \right) dx dy = \iint_{\Omega} (-\Delta u) v dx dy + \langle \gamma_1 u, \gamma_0 v \rangle$$

for $u \in L^2(-\Delta)$ and $v \in H^1(\Omega)$. The duality pairing $\langle \gamma_1 u, \gamma_0 v \rangle$ on $S^* \times S$ extends

$$\int_{\partial\Omega} \frac{du}{dn} v d\sigma$$

in the classical Green's formula (2.7), and we take $\gamma_1 u$ as the generalized normal derivative of $u \in L^2(-\Delta)$.

2.3. Sinc Function Approximation

Most numerical approximation methods, such as interpolation, quadrature and the finite element method, are based on exact relationships that polynomials satisfy. These schemes generally do very well in a region where the function to be approximated is analytic, and very poorly in a neighborhood of a singularity of the function.

The numerical approximation schemes reported in this section have roughly the same accuracy whether or not the function to be approximated has a singularity in its domain. In the absence of singularities, this accuracy is usually not as good as that which may be obtained by polynomial methods, but if singularities are present, this accuracy is much better than that obtained by polynomial methods [27]. All of the approximation methods of the present section are based on exact relationships satisfied by Whittaker's cardinal function $C(f,h)$.

Let f be a function defined on $\mathbb{R}$ and $h > 0$ a stepsize. Then the Whittaker cardinal function is defined by

$$(2.9) \quad C(f,h) = \sum_{k=-\infty}^{\infty} f(kh)S(k,h)$$

whenever this series converges and where

$$(2.10) \quad S(k,h)(x) = \frac{\sin(\pi(x - kh)/h)}{\pi(x - kh)/h}.$$

The function $S(k,h)$ in (2.10) is known as the sinc function based at kh and is defined by continuity at $x = kh$, that is, $S(k,h)(kh) = 1$.

The function $C(f,h)$ was discovered by E. T. Whittaker [33] who studied its mathematical properties and used it as a means of obtaining

alternate expressions of entire functions. He called $C(f,h)$ "a function of royal blood in the family of entire functions whose distinguished properties separate it from its bourgeois brethren." The function $C(f,h)$ has played a role in engineering applications and engineers have referred to (2.9) as the "band limited" or "sinc function" expansion of f .

The identity

$$(2.11) \quad S(k,h)(jh) = \delta_{kj} = \begin{cases} 1 & \text{if } k = j \\ 0 & \text{if } k \neq j \end{cases}$$

is an elementary consequence of (2.10) and, in particular, shows that $C(f,h)(jh) = f(jh)$ for all j . That is, the function $C(f,h)$ provides a formula that interpolates the function f on the infinite set $\{jh\}_{j=-\infty}^{\infty}$. This is only a first hint of some of the beautiful approximation properties of the Whittaker cardinal function. Indeed, an elementary integration shows that the sinc functions satisfy the orthogonality relation

$$\int_{\mathbb{R}} S(j,h)(x) S(k,h)(x) dx = h \delta_{jk}$$

which, when combined with (2.9) shows that

$$\int_{\mathbb{R}} C(f,h)(x) S(k,h)(x) dx = h f(kh).$$

Also, using the fact that the Fourier transform of the characteristic

function on $[-1,1]$ is $(\sin x)/x$ it follows, upon integrating (2.9), that

$$(2.12) \quad \int_R C(f,h)(x) dx = h \sum_{k=-\infty}^{\infty} f(kh).$$

For functions f whose cardinal expansion represent them, it should be pointed out that (2.12) is the familiar trapezoidal quadrature rule. The class of functions for which this is true is characterized in the following definition and theorem.

Definition 2.3. Given $h > 0$, let $B(h)$ be the family of all functions f defined on the complex plane $\mathbb{C}$ such that f is entire, $f \in L^2(\mathbb{R})$, and

$$|f(z)| \leq C \exp(\pi|y|/h),$$

for all $z = x + iy \in \mathbb{C}$ and for some positive constant C .

The following theorem [15] indicates the important role the functions $S(k,h)$ play in the family $B(h)$.

Theorem 2.4. If $f \in B(h)$, then

$$f = C(f,h) = \sum_{k=-\infty}^{\infty} f(kh)S(k,h),$$

$$\int_R f(x) dx = h \sum_{k=-\infty}^{\infty} f(kh)$$

and

$$\int_R (f(x))^2 dx = h \sum_{k=-\infty}^{\infty} (f(kh))^2.$$

Hence, the sequence $\{h^{-1/2}S(k,h)\}$ is a complete orthonormal sequence in $B(h)$.

The derivatives of f can also be expressed as a sinc function expansion [14]. In this direction, set

$$(2.13) \quad \delta_{jk}^{(n)} = S(j,1)^{(n)}(k).$$

In particular,

$$\delta_{jk}^{(0)} = \begin{cases} 1 & \text{if } j = k \\ 0 & \text{if } j \neq k \end{cases}$$

$$\delta_{jk}^{(1)} = \begin{cases} 0 & \text{if } j = k \\ \frac{(-1)^{k-j}}{k-j} & \text{if } j \neq k \end{cases}$$

and

$$\delta_{jk}^{(2)} = \begin{cases} \frac{-\pi^2}{3} & \text{if } j = k \\ \frac{-2(-1)^{k-j}}{(k-j)^2} & \text{if } j \neq k. \end{cases}$$

Theorem 2.5. If $f \in B(h)$, then $f' \in B(h)$ and

$$f^{(n)}(kh) = h^{-n} \sum_{j=-\infty}^{\infty} \delta_{jk}^{(n)} f(jh)$$

and therefore,

$$f^{(n)} = h^{-n} \sum_{k=-\infty}^{\infty} \left(\sum_{j=-\infty}^{\infty} \delta_{jk}^{(n)} f(jh) \right) S(k, h).$$

A class of analytic functions for which the identities in Theorems 2.4 and 2.5 are no longer exact, but form extremely accurate approximations is summarized in

Definition 2.6. For $d > 0$, let $B(D_d)$ denote the family of functions f such that f is analytic in the infinite strip

$$D_d = \{x + iy : |y| < d\},$$

$$\int_{-d}^d |f(x + iy)| \, dy \longrightarrow 0 \quad \text{as } x \longrightarrow \pm \infty,$$

and

$$N(f, d) = \lim_{y \rightarrow d^-} \left\{ \int_{\mathbb{R}} |f(x + iy)| \, dx + \int_{\mathbb{R}} |f(x - iy)| \, dx \right\} < \infty.$$

The uniform error of the approximation of a function from $B(D_d)$ has been calculated by Stenger [25].

Theorem 2.7. Assume that $f \in B(D_d)$. Then for all $h > 0$

$$\|f - C(f, h)\|_{\infty} \leq \frac{N(f, d)}{2\pi d \sinh(\pi d/h)}.$$

Lundin and Stenger [14] have extended the above result to the approximation of the derivatives of f by the derivatives

$$C(f, h)^{(n)}(x) = \frac{d^n}{dx^n} C(f, h)(x)$$

of the cardinal function.

Theorem 2.8. Assume that $f \in B(D_d)$ and suppose $\pi d/h \geq 1$. Then for $n \geq 1$

$$\|f^{(n)} - C(f, h)^{(n)}\|_{\infty} \leq \frac{n! e N(f, d)}{2\pi d \sinh(\pi d/h)} \left(\frac{\pi}{h}\right)^n.$$

In practice, a truncation of the expansion $C(f, h)$ (or $C(f, h)^{(n)}$) is used to approximate f (or $f^{(n)}$). In this direction, define

$$(2.14) \quad C_N(f, h) = \sum_{k=-N}^N f(kh) S(k, h)$$

and

$$(2.15) \quad C_N(f, h)^{(n)} = \sum_{k=-N}^N f(kh) S(k, h)^{(n)}.$$

These approximations are quite accurate for $f \in B(D_d)$, but without further assumptions on the growth of f the N in (2.14) or (2.15) may be inordinately large for economical approximations. A convenient growth assumption for f is given by the inequality

$$(2.16) \quad |f(x)| \leq C e^{-\alpha|x|} \quad \text{for } x \in \mathbb{R},$$

where C and α are positive constants. The following theorem [14] bounds the uniform error in the truncation of the cardinal expansion of f and the truncation of the cardinal expansion of $f^{(n)}$.

Theorem 2.9. Assume that $f \in B(D_d)$ with

$$|f(x)| \leq C e^{-\alpha|x|} \quad \text{for } x \in \mathbb{R},$$

where C and α are positive constants. Then for all $h > 0$

$$(2.17) \quad \|f - c_N(f, h)\|_\infty \leq \frac{N(f, d)}{2\pi d \sinh(\pi d/h)} + \frac{2C}{\alpha h} e^{-\alpha N h}$$

and for $\pi d/h > 1$ and $n \geq 1$

$$(2.18) \quad \|f^{(n)} - c_N^{(n)}(f, h)\|_\infty \leq \frac{n! e N(f, d)}{2\pi d \sinh(\pi d/h)} \left(\frac{\pi}{h}\right)^n + \frac{2C}{(n+1)\alpha\pi} \left(\frac{\pi}{h}\right)^{n+1} e^{-\alpha N h}.$$

Notice that the selection $h = (\pi d/\alpha N)^{1/2}$ in (2.17) and (2.18) imply that the error terms are bounded by

$$K N^{(n+1)/2} \exp\{-(\pi d \alpha N)^{1/2}\}, \quad \text{for } n \geq 0.$$

The positive constant K depends only on n, d, f and α .

The extension of some of the above formulas to finite and semi-infinite intervals, and more generally, to contours may be carried out via the theory of conformal mappings [27]. In this direction the following definition is fundamental.

Definition 2.10. Let D be a simply connected domain, with boundary ∂D , let a and b ($a \neq b$) be points of ∂D , and let D_d be defined as in definition 2.6. Let ϕ be a conformal map of D onto D_d , such that $\phi(a) = -\infty$ and $\phi(b) = \infty$. Let $\psi = \phi^{-1}$ denote the inverse map and set

$$\Gamma = \{ \psi(x) : -\infty \leq x \leq \infty \}$$

(see Figure 2.1). Given ϕ and ψ , denote by $z_k = z_k(h)$ the points

$$z_k = \psi(kh) \quad \text{for } k = 0, \pm 1, \pm 2, \dots$$

Let $B(D)$ denote the family of all functions F that are analytic in D such that

$$\int_{\psi(u+L)} |F(z) dz| \longrightarrow 0 \quad \text{as } u \longrightarrow \pm\infty,$$

where

$$L = \{iy : |y| \leq d\},$$

and such that

$$N(F, D) = \lim_{C_1 \rightarrow \partial D} \inf_{C_1 \subset D} \int_{C_1} |F(z) dz| < \infty.$$

If $F \in B(D)$, then f defined by

$$f = \psi' [F \circ \psi]$$

is in $B(D_d)$, where $B(D_d)$ is defined in Definition 2.6. Then for all

$z \in D$, the approximation

$$F \cong \sum_{k=-N}^N F(z_k) S(k, h) \circ \phi$$

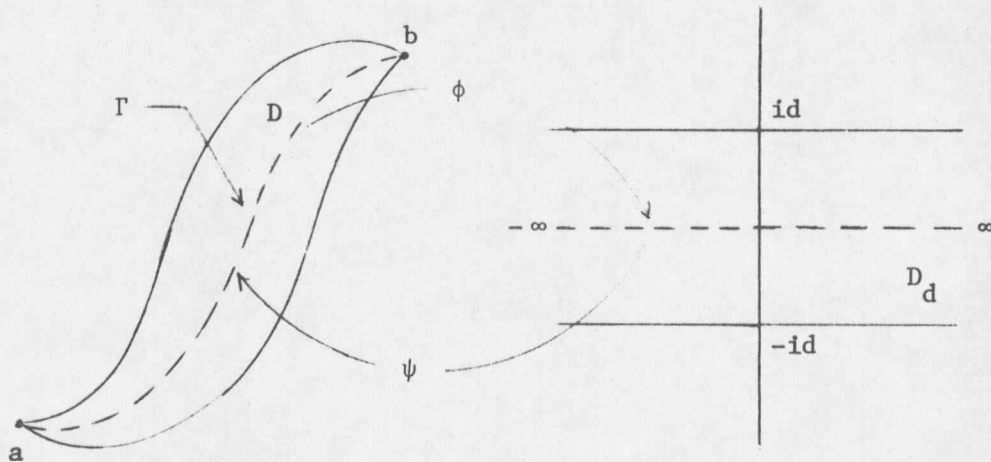


Figure 2.1. The domains D and D_d .

is the analogue of (2.9) in the above. If F also satisfies

$$(2.19) \quad |F(z)| \leq C e^{-\alpha|\phi(z)|} \quad \text{for } z \in \Gamma,$$

where C and α are positive constants, then f satisfies (2.16) and the truncated cardinal expansion

$$F \cong \sum_{k=-N}^N F(z_k) S(k, h) \circ \phi$$

satisfies error estimates of the form (2.17).

For example, to approximate a function F on $\Gamma = [0, \infty)$, define the domain

$$(2.20) \quad D = \{z : |\arg z| < d\}.$$

The sector D is shown in Figure 2.2.

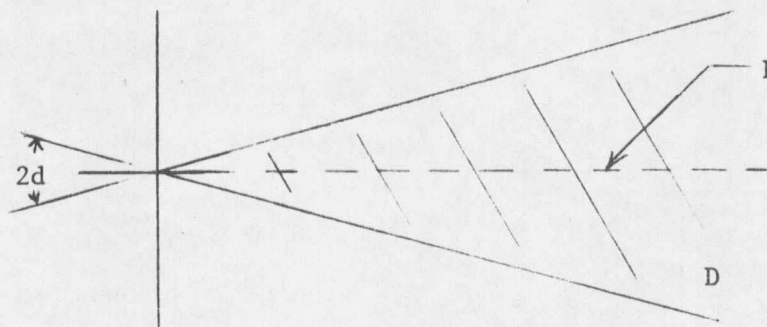


Figure 2.2. The domain $D = \{z : |\arg z| < d\}$.

The functions ϕ and ψ and the points z_k of Definition 2.10 are given by

$$(2.21) \quad w = \phi(z) = \log z,$$

$$z = \psi(w) = e^w,$$

and

$$(2.22) \quad z_k = e^{kh} \quad \text{for } k = 0, \pm 1, \pm 2, \dots.$$

Assume again that $\Gamma = [0, \infty)$ and define, for $0 < d \leq \frac{\pi}{2}$, the domain

$$(2.23) \quad D = \{z : |\arg \sinh z| < d\}.$$

The pointed tubular region D is shown in Figure 2.3.

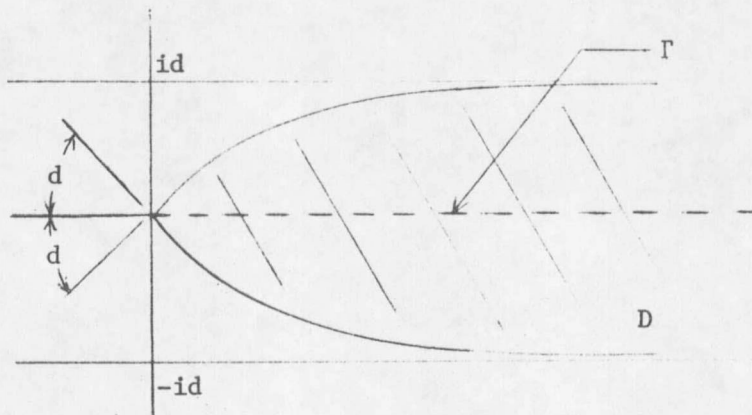


Figure 2.3. The domain $D = \{z : |\arg \sinh z| < d\}$.

The functions ϕ and ψ and the points z_k of Definition 2.10 are given by

$$(2.24) \quad w = \phi(z) = \log \sinh z,$$

$$z = \psi(w) = \log (e^w + (e^{2w} + 1)^{1/2}),$$

and

$$(2.25) \quad z_k = \log (e^{kh} + (e^{2kh} + 1)^{1/2}) \quad \text{for } k = 0, \pm 1, \pm 2, \dots$$

The importance of the conformal maps ϕ in (2.21) and (2.24) in terms of approximate integration is contained in the following theorem [27].

Theorem 2.11. Assume that $f \in B(D)$ where D is given by (2.20) or (2.23). If f satisfies (2.19) for all $x \in \Gamma = [0, \infty)$, then by taking $h = (2\pi d/\alpha N)^{1/2}$ it follows that

$$(2.26) \quad \left| \int_0^\infty f(x) dx - h \sum_{k=-N}^N \frac{1}{\phi'(z_k)} f(z_k) \right| \leq C e^{-(2\pi d \alpha N)},$$

where the z_k are defined as in (2.22) or (2.25) and C is a positive constant depending only on f , α and d .

The following approximations can be obtained by an application of

Theorem 2.11 [27]. In each case, the function f is replaced by its cardinal expansion

$$f \approx \sum_{k=-N}^N f_k S(k, h) \circ \phi,$$

where $f_k = f(z_k)$, and the approximation (2.26) is applied in combination with Equations (2.11) and (2.13). The formulas

$$(2.27) \quad \int_0^{\infty} \frac{g(x)f(x)}{\phi'(x)} S(j, h) \circ \phi(x) \, dx \approx \frac{hg_j f_j}{(\phi'_j)^2},$$

$$(2.28) \quad \int_0^{\infty} \frac{g(x)f'(x)}{\phi'(x)} S(j, h) \circ \phi(x) \, dx \approx \frac{g_j}{\phi'_j} \sum_{k=-N}^N f_k \delta_{kj}^{(1)},$$

and

$$(2.29) \quad \int_0^{\infty} \frac{f''(x)}{\phi'(x)} S(j, h) \circ \phi(x) \, dx \approx \sum_{k=-N}^N f_k \left(\frac{\phi''_j}{(\phi'_j)^2} \delta_{kj}^{(1)} + \frac{1}{h} \delta_{kj}^{(2)} \right)$$

will be used to show that the sinc-collocation method of Chapter 3 is indeed a Galerkin scheme as described by equations (1.2) and (1.3).

To facilitate the description of the sinc-collocation method and

the convergence argument in Chapter 3, define two $(2N + 1) \times (2N + 1)$ matrices by

$$I_N^{(1)} = [\delta_{ij}^{(1)}] = \begin{bmatrix} 0 & -1 & \frac{1}{2} & \frac{-1}{3} & \cdots & \frac{1}{2N} \\ 1 & 0 & -1 & \frac{1}{2} & \cdots & \frac{-1}{2N-1} \\ & & \dots & \dots & \dots & \dots \\ \frac{-1}{2N} & \frac{1}{2N-1} & \frac{-1}{2N-2} & \frac{1}{2N-3} & \cdots & 0 \end{bmatrix}$$

and

$$I_N^{(2)} = [\delta_{ij}^{(2)}] = \begin{bmatrix} \frac{-\pi^2}{3} & \frac{2}{1^2} & \frac{-2}{2^2} & \frac{2}{3^2} & \cdots & \frac{-2}{(2N)^2} \\ \frac{2}{1^2} & \frac{-\pi^2}{3} & \frac{2}{1^2} & \frac{-2}{2^2} & \cdots & \frac{2}{(2N-1)^2} \\ & & \dots & \dots & \dots & \dots \\ \frac{-2}{(2N)^2} & \frac{2}{(2N-1)^2} & \frac{-2}{(2N-2)^2} & \frac{2}{(2N-3)^2} & \cdots & \frac{-\pi^2}{3} \end{bmatrix}$$

The following theorem can be found in [26].

Theorem 2.12. $I_N^{(1)}$ is a skew-symmetric matrix having eigenvalues $i\lambda_k^{(1)}$, where

$$-\pi < \lambda_k^{(1)} < \pi,$$

and $I_N^{(2)}$ is a symmetric matrix having eigenvalues $-\lambda_k^{(2)}$, where

$$4 \sin^2 [\pi/(4N + 4)] < \lambda_k^{(2)} < \pi^2.$$

Corollary 2.13. $\|I_N^{(1)}\|_2 < \pi$.

Proof. Recall that the matrix L^2 -norm $\|A\|_2$ and the spectral radius $\rho(A)$ of a square matrix A satisfy,

$$[\rho(A)]^2 = \rho(A^2) \text{ and } \|A\|_2^2 = \rho(A^T A).$$

Therefore, using the result of Theorem 2.12,

$$\begin{aligned} \|I_N^{(1)}\|_2^2 &= \rho((I_N^{(1)})^T I_N^{(1)}) = \rho(-(I_N^{(1)})^2) \\ &= \rho((I_N^{(1)})^2) = (\rho(I_N^{(1)}))^2 < \pi^2. \end{aligned}$$

The above are but a few of the many approximation properties of the Whittaker cardinal function. These properties have been the basis of numerous accurate numerical processes, Lund's numerical evaluation

of integral transforms [13], Schwing's approximation of solutions of integral equations [21], and Stenger's approximation of solutions of boundary value problems [26], to name just a few. Although the sinc function techniques have been used for many problems, they have not been applied to the numerically complex question of estimating the eigenvalues of a differential or integral equation. The work in Chapter 3 is a step in this direction. The error bounds given above for the approximation of a function and its derivative are used to prove the convergence of the eigenvalues of the discrete sinc-collocation method to the true eigenvalues of the continuous problem.

2.4 Perturbed Matrices

In this section, two theorems on perturbed matrices, found in the work of Steward [28], are stated. These theorems along with a corollary are used in the convergence argument of Chapter 3.

For a square matrix A , let $\sigma(A)$ and $\rho(A)$ denote the spectrum and the spectral radius of A , respectively. It is well known that if A is a symmetric matrix, then $\|A\|_2 = \rho(A)$. Using the following theorem, a similar result is obtained for a perturbed symmetric matrix.

Theorem 2.14. Let λ_1 be a simple eigenvalue of the symmetric matrix A and set $\delta = \min \{|\lambda_1 - \lambda| : \lambda \in \sigma(A) \setminus \{\lambda_1\}\}$. Given a matrix E with $\|E\|_2 = \varepsilon$, if ε is sufficiently small, then there is an

eigenvalue Λ of $A + E$ such that

$$|\Lambda - \lambda_1| \leq \varepsilon + (\|A\|_2 + \frac{1}{\delta})O(\varepsilon^2).$$

Corollary 2.15. Suppose the $n \times n$ symmetric matrix A has n distinct eigenvalues. Let E be an $n \times n$ matrix with $\|E\|_2 = \varepsilon$. If ε is sufficiently small, then

$$\rho(A) \leq \rho(A + E) + \varepsilon + (\|A\|_2 + \frac{1}{\delta})O(\varepsilon^2)$$

where $\delta = \min \{|\lambda_1 - \lambda_2| : \lambda_1, \lambda_2 \in \sigma(A), \lambda_1 \neq \lambda_2\}$.

Proof. Suppose ε is small enough so that Theorem 2.14 holds for each $\lambda \in \sigma(A)$. Then given a $\lambda \in \sigma(A)$ there is a $\Lambda(\lambda) \in \sigma(A + E)$ such that

$$|\Lambda(\lambda) - \lambda| \leq \varepsilon + (\|A\|_2 + \frac{1}{\delta})O(\varepsilon^2).$$

Hence,

$$|\lambda| \leq |\Lambda(\lambda)| + |\lambda - \Lambda(\lambda)| \leq |\Lambda(\lambda)| + \varepsilon + (\|A\|_2 + \frac{1}{\delta})O(\varepsilon^2).$$

Now, calculating the spectral radius of A gives

$$\rho(A) = \max \{|\lambda| : \lambda \in \sigma(A)\}$$

$$\leq \max \{|\Lambda(\lambda)| : \lambda \in \sigma(A)\} + \varepsilon + (\|A\|_2 + \frac{1}{\delta})O(\varepsilon^2)$$

$$\leq \max \{|\Lambda| : \Lambda \in \sigma(A + E)\} + \varepsilon + (\|A\|_2 + \frac{1}{\delta})O(\varepsilon^2)$$

$$= \rho(A + E) + \varepsilon + (\|A\|_2 + \frac{1}{\delta})O(\varepsilon^2).$$

Corollary 2.16. Suppose the $n \times n$ symmetric matrix A has n distinct eigenvalues. Let E be an $n \times n$ matrix with $\|E\|_2 = \varepsilon$. If ε is sufficiently small, then

$$(2.30) \quad \|A + E\|_2 = \rho(A + E) + K\varepsilon + (\|A\|_2 + \frac{1}{\delta})O(\varepsilon^2),$$

where $0 \leq K \leq 2$.

Proof. Suppose ε is small enough so that Corollary 2.15 holds.

Recall that for any matrix norm $\|\cdot\|$, $\|B\| \geq \rho(B)$. Hence,

$$\rho(A + E) \leq \|A + E\|_2 \leq \|A\|_2 + \|E\|_2 = \rho(A) + \varepsilon$$

$$\leq \rho(A + E) + 2\varepsilon + (\|A\|_2 + \frac{1}{\delta})O(\varepsilon^2).$$

Now, (2.30) follows immediately.

The second theorem on perturbed matrices concerns the inverse of such a matrix.

Theorem 2.17. Suppose A is a nonsingular matrix.

If $\|A^{-1}E\|_2 < 1$, then $A + E$ is also nonsingular and

$$(A + E)^{-1} = A^{-1} + FA^{-1},$$

where

$$\|F\|_2 \leq \frac{\|A^{-1}E\|_2}{1 - \|A^{-1}E\|_2}.$$

3. THE SINC-COLLOCATION METHOD

3.1 Introduction

This chapter considers the numerical approximation of the eigenvalues of the radial Schrodinger equation

$$(3.1) \quad -u''(x) + q(x)u(x) = \lambda p(x)u(x) \quad \text{for } x \in \mathbb{R}_+ = (0, \infty)$$

subject to the boundary conditions

$$(3.2) \quad u(0) = 0 \text{ and } \lim_{x \rightarrow \infty} u(x) = 0.$$

In most physical problems, the functions q and p are of the general form

$$q(x) = r(x) + s(x)/x^2$$

and

$$p(x) = t(x) + w(x)/x^2,$$

where r , s , t and w are bounded on every finite interval $[0, b]$. It will also be assumed that the functions r , s , t and w are nonnegative on $\mathbb{R}_+$ so that $q(x) \geq 0$ and $p(x) \geq 0$ on $\mathbb{R}_+$.

The Schrodinger problem is an example of a singular Sturm-Liouville system [6]. Unlike a regular Sturm-Liouville system, which always has a discrete spectrum with square integrable eigenfunctions,

the Schrodinger problem can have a discrete spectrum with square integrable eigenfunctions, a continuous spectrum with eigenfunctions which are not square integrable, or a combination of both situations. In the case of the Schrodinger problem, it is the discrete eigenvalues which are sought. Physically, the eigenvalues correspond to the energy levels that the system can assume.

As the method of Section 3.3 is compared to the Rayleigh-Ritz method, Section 3.2 is devoted to a brief discussion of the latter for Equation (3.1) subject to homogeneous boundary conditions on a finite interval. The classical finite element procedure is then introduced for this problem and the adaption of the procedure, by Shore [22-24], is considered.

The sinc-collocation method is developed and its connection with the Galerkin procedure, as described by (1.2) and (1.3), is identified in Section 3.3. After the convergence of the approximate eigenvalues is established in Section 3.4, the method is applied to several test examples in the final Section 3.5.

3.2 Finite Element Approximation.

Consider the Sturm-Liouville eigenvalue problem

$$(3.3) \quad -u''(x) + q(x)u(x) = \lambda p(x)u(x) \quad \text{for } x \in I = [0, b],$$

where q and p are bounded positive functions on the interval I and the solution u satisfies

$$(3.4) \quad u(0) = u(b) = 0.$$

The Rayleigh-Ritz method, as opposed to classical finite differences, is not applied directly to the differential Equation (3.3). Instead, an equivalent variational form of (3.3) must first be obtained. In this direction, a generalized integration by parts formula is derived by considering the weak L^2 -differential operator D^2 (recall Section 2.2). Suppose u is an element of the Sobolev space $H_0^1(I)$ such that $D^2u \in L^2(I)$, then by thinking of D^2u as the weak L^2 -derivative of Du , one obtains

$$(3.5) \quad \int_I (-D^2u)\phi \, dx = \int_I Du D\phi \, dx \quad \text{for all } \phi \in C_0^\infty(I).$$

Now, since $C_0^\infty(I)$ is dense in $H_0^1(I)$ (Section 2.2), (3.5) persists for all $\phi \in H_0^1(I)$. Equation (3.5) is the required integration by parts formula. Suppose that the solution u of (3.3) and (3.4) is an element of $H_0^1(I)$ with a weak L^2 -derivative D^2u . Multiplying both sides of Equation (3.3) by a $v \in H_0^1(I)$ and integrating over I leads, after integrating by parts, to the identity

$$(3.6) \quad \int_I (DuDv + quv) \, dx = \lambda \int_I puv \, dx.$$

Equation (3.6) is usually called the weak or variational form of the eigenvalue problem (3.3) and (3.4). Notice that if a number λ and a function $u \in H_0^1(I)$ satisfy (3.6) for all $v \in H_0^1(I)$, then the pair, λ and u , is a solution of (3.3) and (3.4) where the derivatives in (3.3) are taken as weak L^2 -derivatives and the equality is in the sense of $L^2(I)$ equality. Indeed, for any $v \in C_0^\infty(I)$, Equation (3.6) can be rewritten as the following statement about distributions on $C_0^\infty(I)$ (recall Section 2.2).

$$(3.7) \quad (-DF_{Du} + F_{qu})(v) = F_{\lambda pu}(v)$$

where F_w is the distribution induced by the L^2 -function w . Hence, the distribution $-DF_{Du}$ is equal to the distribution $F_{(\lambda pu - qu)}$, and Du has a weak L^2 -derivative D^2u . Equation (3.7) may now be written as

$$\int_I (-D^2u + qu)v \, dx = \int_I \lambda puv \, dx \quad \text{for all } v \in C_0^\infty(I).$$

Since $C_0^\infty(I)$ is dense in $L^2(I)$, it follows that

$$-D^2u + qu = \lambda pu$$

as L^2 -functions.

It is the above observation which motivates the Rayleigh-Ritz technique. The eigenvalues and eigenfunctions of (3.3) and (3.4) are approximated by finding the numbers λ_n and functions $u_n \in S_n$ such that

$$(3.8) \quad \int_I (Du_n Dv + qu_n v) dx = \lambda_n \int_I p u_n v dx \quad \text{for all } v \in S_n,$$

where S_n is a finite dimensional subspace of $H_0^1(I)$.

If S_n has the basis $\{\phi_1, \phi_2, \dots, \phi_n\}$, then a solution of the above problem is expressed in terms of the trial functions

$$u_n = \sum_{k=1}^n a_k \phi_k$$

and (3.8) reduces to the discrete system

$$(3.9) \quad \sum_{k=1}^n a_k \int_I (D\phi_k D\phi_j + q\phi_k \phi_j) dx = \lambda_n \sum_{k=1}^n a_k \int_I p\phi_k \phi_j dx$$

for $j = 1, 2, \dots, n$, for the unknown coefficients $\{a_k\}_{k=1}^n$. Notice that (3.9) is the orthogonality of the error in the general Galerkin scheme (1.3).

Equation (3.9) is a system of linear algebraic equations which can be written in the matrix form

$$(3.10) \quad Aa = \lambda_n Ba$$

where A is the $n \times n$ symmetric matrix with (k, j) entry

$$\int_I (D\phi_k D\phi_j + q\phi_k \phi_j) dx,$$

the $n \times n$ symmetric matrix B has (k, j) entry

$$\int_I p\phi_k \phi_j dx,$$

and a is the column vector $a = [a_1, a_2, \dots, a_n]^T$.

Standard techniques are now employed to calculate the eigenvalues of (3.10). These eigenvalues are the approximations to some of the eigenvalues of the continuous problem (3.3) and (3.4) [29]. The finite element method is exactly the Ritz technique with a special choice of basis elements ϕ_k . The interval I is divided into $n - 1$ subintervals defined by the partition $\pi : 0 = x_1 < x_2 < \dots < x_n = b$. The ϕ_k are piecewise polynomials, often cubic splines, defined in such a way that each ϕ_k is nonzero only on a few of the subintervals. The matrices A and B in (3.10) are therefore sparse systems, and theoretically, large systems can be handled.

Birkhoff et al. [5] obtained error bounds for the approximate eigenvalues obtained via the finite element procedure, with cubic spline basis, for the case when the second derivatives of p and q in (3.3) are piecewise continuous, and their jump discontinuities (if any) occur at the nodes $\{x_i\}_{i=1}^n$ of the partition π . In this situation, the eigenfunctions are at least four times continuously differentiable on $[0, b]$, except for possible finite jumps in $u^{(iv)}$ at the nodes of π .

Since the eigenfunctions of the Schrodinger equation satisfy (3.2), Shore [22-24] suggests replacing the Schrodinger problem by the problem (3.3) and (3.4). Shore gives numerical results for the

Schrodinger problem with a number of potentials (different functions q) using cubic B-spline basis functions.

Schoombie and Botha [20] analyze Shore's method. They obtain error bounds composed of two parts. The first part of the error is introduced by the finite element approximation to the continuous problem. This error is similar to that in [5]. The second portion of the error is introduced at the very outset of Shore's procedure. That is, the replacement of the infinite interval in (3.1) and (3.2) by the finite interval in (3.3) and (3.4). The analysis in [20] addresses the question of the error incurred in choosing the truncation point b in (3.3) and (3.4) assuming a known asymptotic behavior of the eigenfunctions. Knowledge of the latter is useful in the applications of the method of the next section and is more fully discussed in Section 3.5.

3.3 Sinc-Collocation Approximation

To use the sinc function approximations for a function and its derivatives, as developed in Section 2.3, it is convenient to transform the radial Schrodinger problem (3.1) and (3.2) onto the whole real line. Either of the mappings $x = e^t$ or $x = \log(e^t + (e^{2t} + 1)^{1/2})$ effect the required transformation. Indeed, these transformations are the inverses, ψ , of the conformal mappings ϕ defined by (2.21) and (2.24). For simplicity and clarity, the

following development is carried out using the transformation $x = e^t$. The derivations using $x = \log (e^t + (e^{2t} + 1)^{1/2})$ are similar in nature and will be omitted. However, this mapping is applied in the examples in Section 3.5.

Using the change of independent variable $x = e^t$, Equation (3.1) is transformed into

$$(3.11) \quad v''(t) - v'(t) - e^{2t}q(e^t)v(t) = -\lambda e^{2t}p(e^t)v(t) \quad \text{for } t \in \mathbb{R}.$$

The boundary conditions (3.2) take the form

$$(3.12) \quad \lim_{x \rightarrow \pm \infty} v(t) = 0.$$

Now, if λ_0 is an eigenvalue of (3.1) and (3.2) with corresponding eigenfunction u_0 , then the pair, λ_0 and $v_0(t) = u_0(e^t)$, is a solution of (3.11) and (3.12).

A trial solution $v(h, N)$ for the problem (3.11) and (3.12) has the assumed form

$$v(h, N) = \sum_{k=-N}^N v_k S(k, h)$$

with $h = C_0 N^{-1/2}$, where C_0 is a positive constant, and $S(k, h)$ is defined in (2.10). That is, $v(h, N)$ is a cardinal-type expansion of a true eigenfunction. Notice that $v(h, N)$ satisfies the boundary conditions (3.12), and the coefficients v_k are determined by the

requirement that $v(h, N)$ satisfy (3.11) at each of the nodes $t = -Nh, (-N + 1)h, \dots, Nh$. The resulting discrete system takes the form

$$(3.13) \quad \sum_{k=-N}^N v_k S(k, h)''(jh) - \sum_{k=-N}^N v_k S(k, h)'(jh) - e^{2jh} q(e^{jh}) \sum_{k=-N}^N v_k S(k, h)(jh) = -\lambda e^{2jh} p(e^{jh}) \sum_{k=-N}^N v_k S(k, h)(jh)$$

for $j = -N, (-N + 1), \dots, N$.

Using the matrices $I_N^{(1)}$ and $I_N^{(2)}$ introduced at the end of Section 2.3, Equation (3.13) can be written as the matrix eigenvalue-eigenvector problem

$$(3.14) \quad [I_N^{(2)} + hI_N^{(1)} - h^2 D_N(e^{2t} q(e^t))] \vec{v}_N = \lambda [h^2 D_N(-e^{2t} p(e^t))] \vec{v}_N$$

where $\vec{v}_N = [v_{-N}, v_{-N+1}, \dots, v_N]^T$ and $D_N(f(t))$ is the $(2N + 1) \times (2N + 1)$ diagonal matrix with diagonal entries $f(-Nh), f((-N + 1)h), \dots, f(Nh)$. The eigenvalues of (3.14) will be shown to approximate the eigenvalues of the Schrodinger problem (3.11) and (3.12).

It should be pointed out that the system (3.14) to determine the unknown coefficients $\{v_k\}_{k=-N}^N$ is referred to in the literature as a collocation scheme. To relate the collocation procedure to a Galerkin development such as (1.2) and (1.3), consider the space H of square integrable functions over $\mathbb{R}_+$ with respect to the weight function

$\frac{1}{\phi'(t)} = t$, where $\phi(t) = \log t$ is the conformal map of Section 2.3.

That is, H is the Hilbert space of functions with the inner product defined by

$$(u, v) = \int_R u(x)v(x) \frac{1}{\phi'(x)} dx.$$

For an eigenvalue λ of the Schrodinger problem, it is assumed that the differential operator

$$L = -\frac{d^2}{dx^2} + (q(x) - \lambda p(x))$$

is a mapping from H back into H . An eigenfunction u associated with the eigenvalue λ is then approximated by the sinc function expansion

$$u(h, N) = \sum_{k=-N}^N u_k S(k, h) \circ \phi$$

where again $h = C_0 N^{-1/2}$ for some positive constant C_0 . The general Galerkin method (recall Section 1.1) enables one to determine the $u_k \approx u(x_k)$ by solving the linear system of equations

$$(3.15) \quad (L(u(h, v)), S(j, h) \circ \phi) = 0 \quad \text{for } j = -N, (-N + 1), \dots, N.$$

Using the quadrature rules (2.27) and (2.29) in Equation (3.15), it follows [26] that the discrete system (3.14) yields precisely the same solution as the system (3.15). Indeed, for the present method, Galerkin and collocation are synonyms.

3.4 Convergence of the Eigenvalues.

Let λ_0 be an eigenvalue of the transformed problem (3.11) and (3.12) with corresponding eigenfunction v_0 . In this section, it is shown that there is an eigenvalue Λ of the matrix eigenvalue problem (3.14) which approximates λ_0 . Further, the error $|\Lambda - \lambda_0|$ is related to how well

$$(3.16) \quad C_N(v_0, h)''(t) - C_N(v_0, h)'(t) - e^{2t} q(e^t) C_N(v_0, h)(t)$$

approximates

$$(3.17) \quad v_0''(t) - v_0'(t) - e^{2t} q(e^t) v_0(t)$$

at the points $t = -Nh, (-N+1)h, \dots, Nh$, where $h = O(N^{-1/2})$. In (3.16), $C_N(v_0, h)$ is the truncated sinc function expansion of v_0

$$C_N(v_0, h) = \sum_{k=-N}^N v_0(kh) S(k, h).$$

Since the pair λ_0 and v_0 is a solution of Equation (3.11), the following holds for all $k = -N, (-N+1), \dots, N$

$$v_0''(kh) - v_0'(kh) - e^{2kh} q(e^{kh}) v_0(kh) = -\lambda_0 e^{2kh} p(e^{kh}) v_0(kh),$$

which is more conveniently expressed as the matrix equation

$$(3.18) \quad \begin{bmatrix} v_0''(-Nh) - v_0'(-Nh) - e^{-2Nh} q(e^{-Nh}) v_0(-Nh) \\ \vdots \\ v_0''(0) - v_0'(0) - e^0 q(e^0) v_0(0) \\ \vdots \\ v_0''(Nh) - v_0'(Nh) - e^{2Nh} q(e^{Nh}) v_0(Nh) \end{bmatrix} = \lambda_0 D_N(-e^{2t} p(e^t)) \vec{v}_N$$

where $\vec{v}_N = [v_0(-Nh), v_0((-N+1)h), \dots, v_0(Nh)]^T$. Multiplying (3.18) by h^2 and subtracting the quantity

$$[I_N^{(2)} + hI_N^{(1)} - h^2 D_N(e^{2t} q(e^t))] \vec{v}_N$$

from both sides of the resulting expression yields the identity

$$(3.19) \quad h^2 \overrightarrow{\Delta_N v_0} = -[I_N^{(2)} + hI_N^{(1)} - h^2 D_N(e^{2t} q(e^t)) - \lambda_0 h^2 D_N(-e^{2t} p(e^t))] \vec{v}_N$$

where $\overrightarrow{\Delta_N v_0} = [\Delta v_0(-Nh), \Delta v_0((-N+1)h), \dots, \Delta v_0(Nh)]^T$ and $\Delta v_0(kh)$ is the difference of (3.16) and (3.17) at $t = kh$.

Now, if λ_0 is not an eigenvalue of (3.19), then the matrix

$$B_N = [I_N^{(2)} + hI_N^{(1)} - h^2 D_N(e^{2t} q(e^t)) - \lambda_0 h^2 D_N(-e^{2t} p(e^t))]$$

is nonsingular, and (3.19) can be rewritten as

$$h^2 B_N^{-1} \overrightarrow{A_N v_0} = -\vec{v}_N.$$

The inequality

$$(3.20) \quad \|\vec{v}_N\|_2 \leq h^2 \|B_N^{-1}\|_2 \|\overrightarrow{A_N v_0}\|_2$$

is elementary and it is seen that the norm of the inverse of the matrix B_N requires examination. For convenience and to make use of the results of Section (2.4) B_N is expressed as the sum

$$B_N = A_N + E_N$$

where A_N is the symmetric matrix

$$A_N = I_N^{(2)} - h^2 D_N(e^{2t} q(e^t)) - \lambda_0 h^2 D_N(-e^{2t} p(e^t))$$

and E_N is the skew-symmetric matrix

$$E_N = h I_N^{(1)}.$$

To make use of the results in Section 2.4, the following assumption is required.

Assumption 3.1 It is assumed that the functions q and p are such that for each N , A_N has only simple eigenvalues with

$$(3.21) \quad \min \{|\lambda| : \lambda \in \sigma(A_N)\} \geq C_1 > 0$$

and that for some $\beta < 1/2$

$$(3.22) \quad \delta_N = \min \left\{ \left| \frac{1}{\lambda_1} - \frac{1}{\lambda_2} \right| : \lambda_1, \lambda_2 \in \sigma(A_N), \lambda_1 \neq \lambda_2 \right\} \geq C_2 N^{-\beta} > 0,$$

where C_1 and C_2 are constants independent of N .

The inequality (3.21) and the symmetry of A_N imply that

$$\|A_N^{-1}\|_2 \leq \frac{1}{C_1},$$

while the definition of E_N and Corollary 2.13 yield

$$\|E_N\|_2 = h \|I_N^{(1)}\|_2 < h\pi.$$

It therefore follows that

$$\|A_N^{-1}E_N\|_2 \leq \|A_N^{-1}\|_2 \|E_N\|_2 < \frac{h\pi}{C_1} = o(N^{-1/2}).$$

For sufficiently small h , $\|A_N^{-1}E_N\|_2 < 1$ so that an application of Theorem 2.17 yields

$$(A_N + E_N)^{-1} = A_N^{-1} + F_N A_N^{-1}$$

where

$$\|F_N\|_2 \leq \frac{\|A_N^{-1}E_N\|_2}{1 - \|A_N^{-1}E_N\|_2}.$$

The estimate

$$\varepsilon = ||F_N A_N^{-1}||_2 \leq ||F_N|| ||A_N^{-1}||_2 < \frac{h\pi/C_1^2}{1 - \frac{h\pi}{C_1}} = O(N^{-1/2})$$

which when combined with Corollary 2.16 imply that

$$(3.23) \quad ||(A_N + E_N)^{-1}||_2 = \rho((A_N + E_N)^{-1}) + C_3 \varepsilon + (||A_N^{-1}||_2 + \delta_N^{-1})O(\varepsilon^2),$$

where $0 \leq C_3 \leq 2$ and δ_N is given by (3.22). Using each of (3.22) and (3.23) and selecting $h = C_0 N^{-1/2}$ shows that

$$(||A_N^{-1}||_2 + \delta_N^{-1})O(\varepsilon^2) = O(\varepsilon).$$

Combining this equality with (3.23) yields the final estimate

$$(3.24) \quad ||(A_N + E_N)^{-1}||_2 \leq \rho((A_N + E_N)^{-1}) + C_4 \varepsilon$$

where C_4 is a positive constant. Recalling that $B_N = A_N + E_N$ and substituting (3.24) into the right hand side of (3.20) leads to the inequality

$$(3.25) \quad ||\bar{v}_N||_2 \leq h^2 \left(\frac{1}{|\Lambda - \lambda_0|} + C_{4,\varepsilon} \right) ||\bar{\Delta}_N v_0||_2,$$

where Λ is an eigenvalue of (3.14).

Assume now that the eigenfunction v_0 of (3.11) satisfies the hypotheses of Theorem 2.9 (this is true for all of the examples of

Section 3.5). The comments following that theorem and the interpolatory nature of $C_N(v_0, h)$ (the sinc expansion of v_0 in (3.16)) imply that there exist a positive constant C_5 such that

$$|\Delta v_0(kh)| \leq C_5 N^{3/2} e^{-(\pi d a N)^{1/2}} \quad \text{for } k = -N, \dots, N.$$

Therefore,

$$\|\overrightarrow{\Delta_N v_0}\|_2 = \left(\sum_{k=-N}^N |\Delta v_0(kh)|^2 \right)^{1/2} \leq C_6 N^2 e^{-(\pi d a N)^{1/2}}$$

for some positive constant C_6 . Thus, $\|\overrightarrow{\Delta_N v_0}\|_2$ goes to zero as N increases.

Finally, since v_0 is an eigenfunction, $v_0 \neq 0$, so it is reasonable to assume that

$$\|\vec{v}_N\|_2 = \left(\sum_{k=-N}^N |v_0(kh)|^2 \right)^{1/2} > C_7 > 0$$

for all sufficiently large N . In particular, for sufficiently small h

$$C_4 h^2 \varepsilon \|\overrightarrow{\Delta_N v_0}\|_2 < \|\vec{v}_N\|_2,$$

so that the inequality (3.25) can be rewritten as

$$|\lambda - \lambda_0| < \frac{h^2 \|\overrightarrow{\Delta_N v_0}\|_2}{\|\vec{v}_N\|_2 - C_4 h^2 \varepsilon \|\overrightarrow{\Delta_N v_0}\|_2}.$$

In summary, the following theorem has been proven.

Theorem 3.2. Let λ_0 be an eigenvalue of the transformed Schrodinger problem (3.11) and (3.12) with corresponding eigenfunction v_0 . Suppose $v_0 \in B(D_d)$, as defined in Definition 2.6, with

$$|v_0(t)| \leq Ce^{-\alpha|t|} \quad \text{for all } t \in \mathbb{R},$$

where C and α are positive constants. Further, suppose the coefficient functions q and p of Equation (3.11) satisfy Assumption 3.1. Then for sufficiently large N and $h = (\pi d/\alpha N)^{1/2}$, there is an eigenvalue Λ of the matrix eigenvalue problem (3.14) such that

$$(3.27) \quad |\Lambda - \lambda_0| < \frac{C_6(\pi d/\alpha)Ne^{-(\pi d\alpha N)^{1/2}}}{\|\vec{v}_N\|_2 - C_4(\pi d/\alpha)\epsilon Ne^{-(\pi d\alpha N)^{1/2}}},$$

where

$$\|\vec{v}_N\|_2 = \left(\sum_{k=-N}^N |v_0(kh)|^2 \right)^{1/2},$$

$$\epsilon = \frac{(\pi C_0/C_1^2)N^{-1/2}}{1 - \frac{\pi C_0}{C_1}N^{-1/2}},$$

and the C_i are positive constants.

Theorem 3.2 states that some eigenvalue Λ of the discrete problem is a good approximation (in fact, $O(Ne^{-\gamma N^{1/2}})$, $\gamma > 0$) to a given eigenvalue λ of the continuous problem. But as $N \rightarrow \infty$, the

theorem fails to identify which eigenvalue Λ is the closest approximation.

Notice that the matrix eigenvalue problem (3.14) has at most $2N + 1$ distinct eigenvalues. However, many of the problems (3.11) and (3.12) have an infinite discrete spectrum

$$\lambda_1 \leq \lambda_2 \leq \lambda_3 \leq \dots$$

where $\lim_{k \rightarrow \infty} \lambda_k = \infty$. Yet, the inequality (3.16) holds for each λ_k .

This is explained by the fact that the eigenfunction v_k corresponding to the eigenvalue λ_k , oscillates more rapidly with increasing k , or equivalently, its derivatives increase in magnitude with k .

Therefore, the error $\Delta v_k(jh)$, with $h = C_0 N^{-1/2}$ held fixed, also increases in magnitude with k . Thus, it is expected that approximations for the smaller eigenvalues of (3.11) and (3.12) can be obtained with a greater degree of accuracy than can the larger eigenvalues. This expectation is born out by the examples of the next section.

3.5 Numerical Implementation and Examples.

In this section, the sinc-collocation method is applied to various test examples. The examples are chosen both for comparison purposes (relative to existing methods) and to expose the versatility

and power of the method.

The first two examples appear in the paper of Schoombie and Botha [20] and the accuracy of the results herein contained are seen to be competitive with their method. It should also be pointed out that each of the examples were computed using different maps ϕ . Example 3.3 was computed using the map ϕ given by (2.21), while Example 3.4 was computed using the map ϕ given by (2.24). With regards to implementation in changing from one map to another, the only changes required in the numerical computation of (3.14) is a minor change in the diagonal matrix D_N and the multiplication of the matrix $I_N^{(1)}$ by a diagonal matrix. The choice of which map to use on the $[0, \infty)$ domain is based on some asymptotic considerations. Since the solutions of (3.1) and (3.2) usually die off rapidly for large x (of the order $e^{-x^2/2}$ and e^{-x} in the cases of Examples 3.3 and 3.4, respectively) either of the maps (2.21) or (2.24) are, in general, applicable. The selection is based, instead, on the oscillatory behavior of the solutions. The conditions defined in Definition 2.10 have been analyzed in [13] where it is shown that the log sinh map is more suitable than the log map for approximation when the function being approximated is oscillatory. This rule of thumb prevails in the present case. It should be pointed out that if circumstances ever demand computation of the higher eigenvalues (the typical cases of

physical interest seem to demand the lower end of the spectrum), then the log sinh map could play a significant role.

The eigenfunctions of the third example are singular at zero, and therefore, polynomial based numerical schemes will generally perform poorly for this example. In particular, due to the form of the errors (involving higher derivatives of the eigenfunctions), it is not clear that the methods of [20] are even applicable to this example. The final example, the classic Hermite differential equation, is included to numerically strengthen the suspicion created in the comments following Theorem 3.2.

Example 3.3. As a first example, consider the Schrodinger equation with a harmonic oscillator potential

$$-u''(x) + \left(x^2 + \frac{2}{x}\right)u(x) = \lambda u(x) \quad \text{for } x \in \mathbb{R}_+$$

subject to the boundary conditions

$$u(0) = 0 \text{ and } \lim_{x \rightarrow \infty} u(x) = 0.$$

The eigenvalues are given by $\lambda_i = 4i + 1$ for $i = 1, 2, \dots$, and the corresponding eigenfunctions are

$$u_i(x) = x^2 y_i(x) e^{-x^2/2}$$

where the y_i are the polynomial solutions of

$$-xy'' + (2x^2 - 4)y' + (5 - \lambda_i)xy = 0.$$

In the development of the previous sections, a sinc function expansion

$$(3.28) \quad v(h, N) = \sum_{k=-N}^N v_k S(k, h)$$

of the transformed eigenfunction $v(t) = u(\phi^{-1}(t))$, with $h = C_0 N^{-1/2}$ and $\phi(x) = \log x$ is used to obtain approximate eigenvalues. It was observed that strict adherence to this convention is not necessary for the acquisition of accurate approximations to these eigenvalues. The following procedure is used to obtain the results reported in Table 3.1. The mapping $\phi(x) = \log x$ is used in conjunction with (3.28) and starting with $h = 1$, successive approximations are calculated for increasing N . It was observed that the smallest eigenvalues for the cases $N = 1, 2$ and 3 were essentially the same. It was concluded that convergence was obtained for the step size $h = 1$. Notice that with $N = 3$ the discrete problem involves the point evaluations $v(k)$, $k = -3, -2, \dots, 3$, of the eigenfunction v . Hence, it was further concluded that for later approximations no substantial gain was to be made in evaluating $v(t)$ for $t > 3$. Now, to increase the accuracy of the approximations a smaller step size h is selected and the corresponding N is chosen so that $hN = 3$. Table 3.1 contains the true eigenvalues of the continuous problem and the smallest eigenvalues of

Table 3.1. Approximate eigenvalues for the radial Schrodinger equation with a harmonic oscillator potential:

$$-u''(x) + (x^2 + \frac{2}{x^2})u(x) = \lambda u(x) \quad \text{for } x \in \mathbb{R}_+$$

with $u(0) = 0$ and $\lim_{x \rightarrow \infty} u(x) = 0$.

True Eigenvalues Approximations using (3.28)
and the log map (2.21)

λ_i	$h = 1.0$ $N = 3$	$h = .5$ $N = 6$	$h = .25$ $N = 12$	$h = .125$ $N = 24$
5	5.476745	4.992981	5.00632	5.000366
9	8.239481	9.080732	8.973630	9.012607
13	-	12.40204	13.01165	12.99365
17	-	-	16.88937	17.02145
21	-	-	21.42238	21.00994
25	-	-	-	25.00346
29	-	-	-	29.00815
33	-	-	-	33.00545
37	-	-	-	37.00822
41	-	-	-	40.95694
45	-	-	-	44.92854
49	-	-	-	48.75737
53	-	-	-	53.65775

each discrete problem associated with the step sizes 1.0, .5, .25, and .125. Four choices of h were used in this example to demonstrate that as h decreases (while hN remains constant) not only are more accurate approximations to the lower eigenvalues obtained, but also a greater number of eigenvalues of the continuous problem are simultaneously approximated.

It was also observed that the symmetric series (3.28) does not necessarily yield optimum approximations in the lower end of the spectrum. A more versatile series is the nonsymmetric cardinal expansion

$$(3.29) \quad v(h, M, N) = \sum_{k=-M}^N v_k S(k, h).$$

For example, with the selection $h = .25$, $M = 18$, and $N = 6$ in (3.29), the first two eigenvalues of the corresponding discrete system were 5.000014 and 8.999512. Notice that for exactly the same amount of work these approximations are significantly more accurate than those presented in the third column of Table 3.1. Indeed, these approximations are more accurate than the case with $h = .125$ where twice as many basis elements are used. These observations are easily explained by the rapid decay of the eigenfunctions $u_1(x)$. The cubic spline procedure of Schoombie and Botha is based on a uniform partition of the interval $[0, b]$. However, they conjectured that for some

problems a non-uniform partition would be more advantageous than the uniform partition. In particular, a partition which is closely spaced near zero but with progressively larger intervals away from zero. This choice of nodes is automatically accomplished in the sinc-collocation method. Moreover, the nonsymmetric series (3.29) with $M > N$ has the effect of concentrating even more nodes near the origin than the symmetric series (3.28).

Example 3.4. Next, consider the Schrodinger equation with a Wood-Saxon potential

$$-u''(x) + u(x) = \lambda \{1 + \exp[(x-r)/a]\}^{-1} u(x) \quad \text{for } x \in \mathbb{R}_+$$

subject to the boundary conditions

$$u(0) = 0 \text{ and } \lim_{x \rightarrow \infty} u(x) = 0,$$

where $r = 5.086855$ and $a = .929853$. In this problem, an analytical solution does not exist in a closed form, however, various investigators have computed the smallest eigenvalue to be approximately $\lambda_1 = 1.424333$ and that the eigenfunctions act asymptotically as $\exp[-(x-r)]$ [20]. It was observed that the sinc-collocation method with $\phi(x) = \log x$, while providing accurate approximations, did not produce estimates of λ as accurate as those of Schoombie and Botha. A possible explanation for this observation

is that the eigenfunctions have a more oscillatory behavior than those of Example 3.1. This combined with the less rapid (as compared to the previous example) dampening of the eigenfunctions would produce a slow rate of convergence of the collocation method. To accommodate eigenfunctions of the above description the sinc-collocation method should be equipped with the mapping $\phi(x) = \log \sinh x$. Table 3.2 contains the results of this procedure, with the symmetric series (3.30), for the Wood-Saxon example. Notice that columns two and three present the results for the same step size $h = .5$ but with different choices of N . The same procedure as described in the previous example for selecting N was used. However, the approximation for the case $h = .5$ and $N = 18$ was so accurate that instead of decreasing h again, the N was increased slightly and a more accurate approximation to the value of λ was achieved. This author has not been able to find accepted estimates for the other eigenvalues λ_k ($k = 2, 3, 4, 5, 6$), however, the values in Table 3.2 certainly indicate convergence and are therefore assumed to be good estimates for these eigenvalues as well.

Example 3.5. Now, consider the Schrodinger equation

$$-u''(x) + \left(x^2 + \frac{1}{4x^2}\right)u(x) = \lambda u(x) \quad \text{for } x \in \mathbb{R}_+$$

Table 3.2. Approximate eigenvalues for the radial Schrodinger equation with a Wood-Saxon potential:

$$-u''(x) + u(x) = \lambda \{1 + \exp[(x - r)/a]\}^{-1} u(x) \quad \text{for } x \in \mathbb{R}_+$$

with $u(0) = 0$ and $\lim_{x \rightarrow \infty} u(x) = 0$, where $r = 5.086855$ and $a = .929853$.

True Eigenvalues	Approximations using (3.28) and the log sinh map (2.24)		
λ_i	$h = 1.0$ $N = 3$	$h = .5$ $N = 6$	$h = .25$ $N = 12$
1.424333	1.424416	1.424360	1.424332
λ_2	2.445411	2.444823	2.444706
λ_3	3.974748	3.972747	3.972347
λ_4	5.974276	5.994216	5.992959
λ_5	8.260024	8.505146	8.502698
λ_6	11.44842	11.50417	11.49909

subject to the boundary conditions

$$u(0) = 0 \text{ and } \lim_{x \rightarrow \infty} u(x) = 0.$$

The eigenvalues are given by $\lambda_i = 4i$ for $i = 1, 2, \dots$, and the corresponding eigenfunctions are

$$u_i(x) = x^{3/2} y_i(x) e^{-x^2/2}$$

where the y_i are the polynomial solutions of

$$-xy'' + (2x^2 - 3)y' + (4 - \lambda)xy = 0.$$

Notice that the eigenfunctions are singular at zero, and therefore, polynomial based numerical methods will generally perform poorly on this example. Table 3.3 contains the results of the sinc-collocation method equipped with the log map. The use of the nonsymmetric approximation (3.29) resulted in only slight improvement of the eigenvalue estimates.

Example 3.6. As a final example, consider the Hermite differential equation

$$-u''(x) + x^2 u(x) = \lambda u(x) \quad \text{for } x \in \mathbb{R}$$

subject to the boundary conditions

Table 3.3. Approximate eigenvalues for a radial Schrodinger equation with singular eigenfunctions:

$$-u''(x) + (x^2 + \frac{1}{4x^2})u(x) = \lambda u(x) \quad \text{for } x \in \mathbb{R}_+$$

with $u(0) = 0$ and $\lim_{x \rightarrow \infty} u(x) = 0$.

True Eigenvalues λ_i	Approximations using (3.28) and the log map (2.21)		
	$h = 1.0$ $N = 3$	$h = .5$ $N = 6$	$h = .25$ $N = 14$
4	3.970246	3.987289	4.006780
8	8.014940	8.296159	7.994587
12	—	11.06397	12.03366
16	—	—	16.03673
20	—	—	20.03673

$$\lim_{x \rightarrow \pm\infty} u(x) = 0.$$

The eigenvalues are given by $\lambda_i = 2i + 1$, $i = 0, 1, \dots$, and the eigenfunctions are

$$u_i(x) = H_i(x)e^{-x^2/2}$$

where the H_i are the Hermite polynomials. Since no transformation is needed for this example the discrete problem has a different form than those of the earlier examples. Using the sinc-function expansion (3.28) for an eigenfunction, the discrete problem is

$$\left[-\frac{1}{2}I_N^{(2)} + D_N(x^2)\right]\bar{v}_N = \lambda\bar{v}_N.$$

Table 3.4 lists some typical eigenvalue approximations. Notice the dramatic increase in accuracy achieved by a single reduction of step size.

It must be emphasized that the approximations in Tables 3.1 through 3.4 contain the smallest eigenvalues for each particular discrete problem. That is, the smallest eigenvalues of the discrete problem approximate the smallest eigenvalues of the continuous problem and there are no extraneous eigenvalues. This phenomenon is not guaranteed by Theorem 3.2, however, it has been observed for

Table 3.4. Approximate eigenvalues for the Hermite differential equation:

$$-u''(x) + x^2 u(x) = \lambda u(x) \quad \text{for } x \in \mathbb{R}$$

with $\lim_{x \rightarrow \pm\infty} u(x) = 0$.

True Eigenvalues Approximations using (3.28)

λ_i	$h = 1.0$ $N = 5$	$h = .5$ $N = 10$
1	.999609	.999999
3	3.005978	3.000000
5	4.942778	5.000000
7	7.201062	6.999998
9	-	9.000000
11	-	11.000000
13	-	13.000004
15	-	15.00027
17	-	17.00142
19	-	19.00607
21	-	21.02165
23	-	23.06545
25	-	25.16377
27	-	27.36334
29	-	29.65257

other approximation schemes such as the finite difference method. In fact, for the latter method it appears as much folklore as fact that this occurrence is expected whenever a finite difference type method is applied to a continuous eigenvalue problem. The examples herein included, lend strong numerical credence that this folklore has some validity in the present sinc-collocation method as applied to the radial Schrodinger equation.

4. THE ANNULAR STRIP METHOD

4.1 Introduction

A heat transfer problem, which arises because of changes in an aging nuclear reactor, is considered in this chapter. The reactor core is a large graphite block through which there are numerous channels. Long hollow fuel rods are inserted into the core channels, and while the reactor is in operation, water is allowed to flow through the fuel rods as well as through the core channels around the fuel rods. In a physically ideal situation, a straight rod is inserted into a straight channel. Figure 4.1 shows the rod partially inserted into the channel. The shaded areas indicate the regions through which water flows. Figure 4.2 shows a cross sectional view of this ideal situation where the clear annular region represents the fuel rod. There are several layers of cladding which form the two boundaries of this annulus and radioactive material lies between the two sets of cladding. Again, the shaded areas indicate the regions through which water flows. During the reaction time, the temperature of the rod increases tremendously. The water running through the rod is heated to extremely high temperatures and the water that flows through the core channel around the fuel rod aids in the dissipation of the heat in the rod.

As the reactor ages, fissures develop in the graphite core and

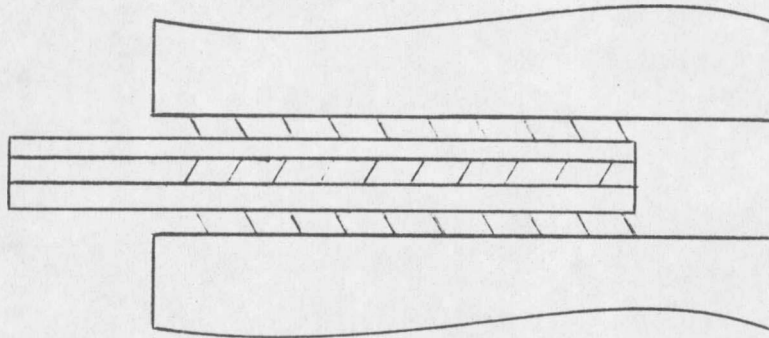


Figure 4.1 Fuel rod partially inserted into core channel.

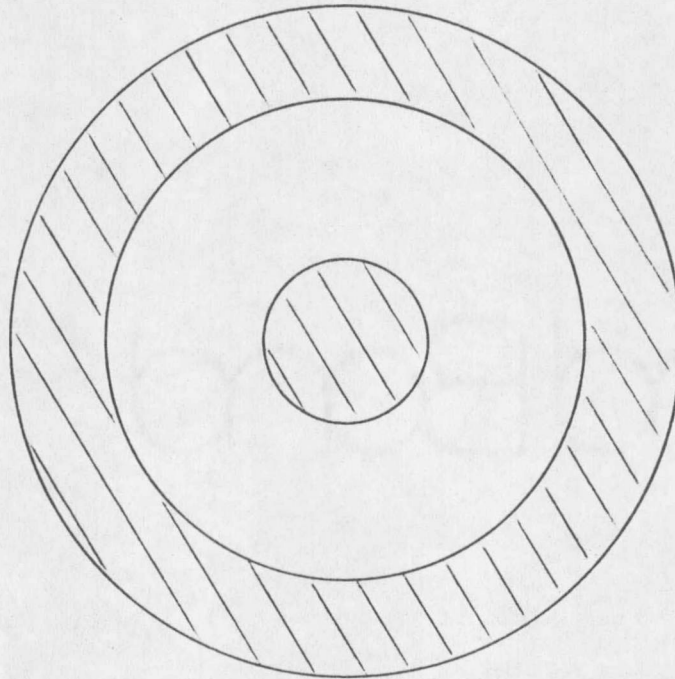


Figure 4.2 Cross section of the fuel rod within the core channel.

the consequent shifting which results cause the core channels to bow. Now, when a straight fuel rod is inserted into the bowed channel, the rod approaches the walls of the channel at certain points. Figure 4.3 shows such a situation and Figure 4.4 is a cross sectional view of the rod and channel at a point where the fuel rod is near the channel wall.

In this new situation, the heat dispersion is not as it was in the physically ideal case. There is now concern that at some point, as the fuel rod approaches the channel wall, the heat dispersion may be extremely poor. So poor, in fact, that the temperature of the rod may increase to a point which will break down the cladding and expose the radioactive material to the water flowing in the core channel.

The problem is to find the temperature within the fuel rod when the rod is at different distances from the channel wall. Once these temperatures are known, they can be used to determine whether or not the reactor with a bowed core channel is safe to keep in operation. The mathematical model of the physical problem is based on a partial differential equation and, after an appropriate transformation, this model is formulated as follows: Find the function $T(x,y)$ defined on the domain Ω (the cross section of the fuel rod), pictured in Figure 4.5, such that

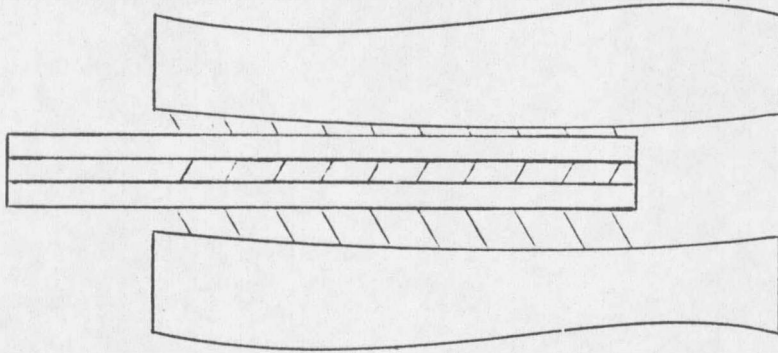


Figure 4.3. Fuel rod partially inserted into a bowed core channel.

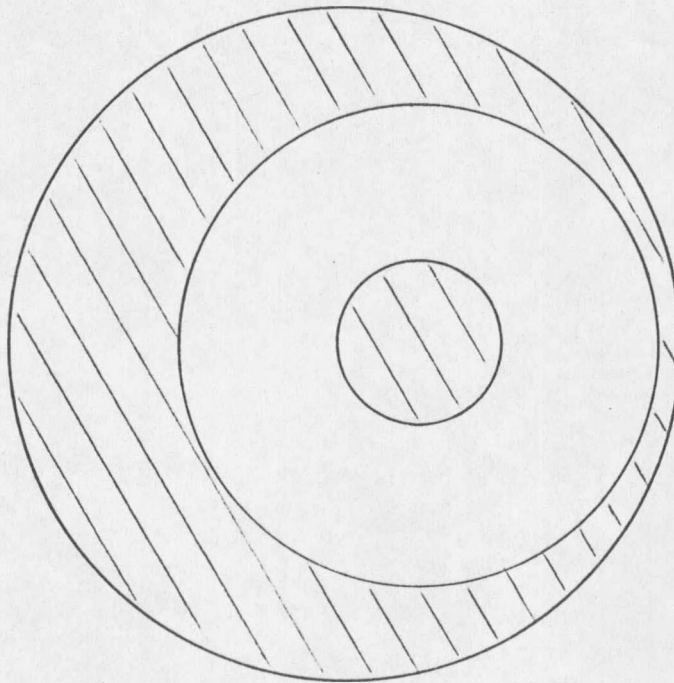


Figure 4.4. Cross section of the fuel rod within the bowed core channel.

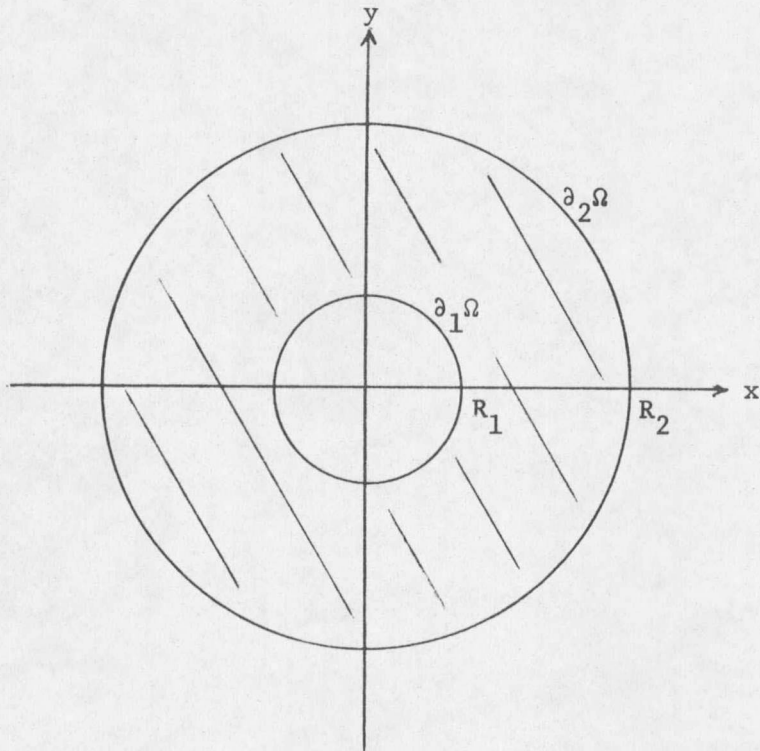


Figure 4.5. The domain Ω with inner boundary $\partial_1\Omega$ and outer boundary $\partial_2\Omega$.

$$(4.1) \quad \Delta T = 0 \quad \text{on } \Omega,$$

$$(4.2) \quad T = 0 \quad \text{on } \partial_1\Omega,$$

and

$$(4.3) \quad \frac{dT}{dn} + gT = k_1T + k_2 \quad \text{on } \partial_2\Omega,$$

where Δ is the Laplacian differential operator

$$\Delta = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2}.$$

$\frac{dT}{dn}$ denotes the derivative of T in the direction of the outward normal to the curve $\partial_2\Omega$, g is a continuous nonnegative function defined on $\partial_2\Omega$, and the k_i ($i = 1, 2$) are constants.

In Section 4.2, the equivalent variational form of the boundary value problem is stated, and the Ritz-Galerkin technique is discussed. The finite element method is described in Section 4.3 and an annular strip approach is developed in Section 4.4. To lend numerical credence to the theory of the method of Section 4.4, the annular strip method is applied to different test problems in Section 4.5.

4.2 The Ritz-Galerkin Technique

A variational form of the boundary value problem (4.1) - (4.3) is required for the implementation of the Ritz-Galerkin technique. In this direction let $C_1^1(\bar{\Omega})$ denote the family of continuously differentiable functions on the closed annulus $\bar{\Omega}$ which are zero on the inner boundary $\partial_1\Omega$. Let $H_1^1(\Omega)$ be the completion of $C_1^1(\bar{\Omega})$ with respect to the Sobolev norm

$$||u||_1 = \left(\iint_{\Omega} (u^2 + (\frac{\partial u}{\partial x})^2 + (\frac{\partial u}{\partial y})^2) dx dy \right)^{1/2}$$

(recall the first construction of a Sobolev space). $H_1^1(\Omega)$ is the subspace of $H^1(\Omega)$ consisting of the functions which are zero on $\partial_1\Omega$. Finally, let $L_1^2(\Delta)$ denote the set of functions $u \in H_1^1(\Omega)$ which have the weak L^2 -derivative Δu . Recall (from the discussion at the end of Section 2.2) that the weak normal derivative γ_1 is defined for functions in $L_1^2(\Delta)$.

With this notation and the assumption on g it is known (see for example Babuska and Aziz [4], Mercier [16], or Oden and Reddy [19]) that the following problems are equivalent.

Problem 4.1. Find the function $T \in L_1^2(\Delta)$ such that

$$\Delta T = 0 \quad \text{in } L^2(\Omega)$$

and

$$\gamma_1 T + g\gamma_0 T = k_1 g + k_2 \quad \text{in } L^2(\partial_2\Omega).$$

Problem 4.2. Find the function $T \in H_1^1(\Omega)$ such that

$$(4.4) \quad \iint_{\Omega} \left(\frac{\partial T}{\partial x} \frac{\partial v}{\partial x} + \frac{\partial T}{\partial y} \frac{\partial v}{\partial y} \right) dx dy + \int_{\partial_2\Omega} g T v \, d\sigma = \int_{\partial_2\Omega} (k_1 g + k_2) v \, d\sigma$$

for all $v \in H_1^1(\Omega)$, where it is understood that the boundary integrals involve the traces of T and v .

In particular, notice that Problem 4.1 is a weak formulation of the original boundary value problem (4.1) - (4.3). However, if T is a classical solution of the original problem, then it is a solution of Problem 4.1, and hence of the variational Problem 4.2.

There is also a partial converse. It will be shown that there is a unique solution of Problem 4.2, and therefore, by the aforementioned equivalence, a unique solution of Problem 4.1. In the case that a classical solution of the original problem (4.1) - (4.3) exists, it must also be the unique solution of Problems 4.1 and 4.2. Thus, the solution of Problem 4.2 must coincide with the classical solution of the original problem if one exists.

The following result, the Lax-Milgram theorem [31], is used to show that Problem 4.2 has a unique solution.

Theorem 4.3. Let H be a Hilbert space, $b(\cdot, \cdot)$ a continuous bilinear form on $H \times H$, and ℓ a continuous linear operator on H . If there is a constant $C > 0$ such that

$$b(u, u) \geq C \|u\|^2 \quad \text{for all } u \in H,$$

then there exists a unique element $T \in H$ such that

$$b(T, v) = \ell(v) \quad \text{for all } v \in H.$$

A bilinear form $b(\cdot, \cdot)$ satisfying the lower bound in the hypotheses of

the Lax-Milgram theorem is called H-elliptic.

Identify the space in the Lax-Milgram theorem as $H = H_1^1(\Omega)$, the bilinear form as

$$(4.5) \quad b(u, v) = \iint_{\Omega} \left(\frac{\partial u}{\partial x} \frac{\partial v}{\partial x} + \frac{\partial u}{\partial y} \frac{\partial v}{\partial y} \right) dx dy + \int_{\partial_2 \Omega} g u v \, d\sigma,$$

and the linear functional as

$$(4.6) \quad \ell(u) = \int_{\partial_2 \Omega} (k_1 g + k_2) u \, d\sigma.$$

Notice that (4.5) and (4.6) are the left and right hand sides of Equation (4.4), respectively. It will be shown that ℓ is continuous and that $b(\cdot, \cdot)$ is continuous and $H_1^1(\Omega)$ -elliptic. The Lax-Milgram theorem then implies the existence of a unique solution to Problem 4.2. Under the assumptions on g , the mapping

$$v \longmapsto \int_{\partial_2 \Omega} (k_1 g + k_2) v \, d\sigma$$

is a continuous linear functional on $L^2(\partial_2 \Omega)$. This, combined with the continuity of the trace operator γ_0 (Section 2.2), implies that ℓ is a continuous linear functional on $H_1^1(\Omega)$.

By a similar argument the mapping

$$(u, v) \longmapsto \int_{\partial_2 \Omega} g u v \, d\sigma$$

is a continuous bilinear form on $H_1^1(\Omega) \times H_1^1(\Omega)$, and it has been shown in Section 2.2 that

$$a(u, v) = \iint_{\Omega} \left(\frac{\partial u}{\partial y} \frac{\partial v}{\partial y} + \frac{\partial u}{\partial x} \frac{\partial v}{\partial x} \right) dx dy$$

is a continuous bilinear form on $H_1^1(\Omega) \times H_1^1(\Omega)$. Hence, the bilinear form (4.5) is continuous on $H_1^1(\Omega) \times H_1^1(\Omega)$. The $H_1^1(\Omega)$ -ellipticity of $b(\cdot, \cdot)$ will require the following lemma.

Lemma 4.4. For any $v \in H_1^1(\Omega)$,

$$(4.7) \quad \iint_{\Omega} v^2 dx dy \leq \frac{(R_2 - R_1)(R_2^2 - R_1^2)}{2R_1} \iint_{\Omega} \left(\left(\frac{\partial v}{\partial x} \right)^2 + \left(\frac{\partial v}{\partial y} \right)^2 \right) dx dy.$$

Proof. It is convenient to use the polar coordinate system in this argument. The annulus Ω in these coordinates is given by

$$\Omega = \{(r, \theta) : R_1 < r < R_2, -\pi < \theta \leq \pi\}.$$

The result (4.7) will first be obtained for any $v \in C_1^1(\Omega)$.

Notice that for any $(r, \theta) \in \Omega$

$$v(r, \theta) = \int_{R_1}^r \frac{\partial v}{\partial r}(t, \theta) dt,$$

and an application of the Cauchy-Schwartz inequality yields

$$(4.8) \quad |v(r, \theta)| \leq \left(\int_{R_1}^r 1^2 dt \right)^{1/2} \left(\int_{R_1}^r \left(\frac{\partial v}{\partial r}(t, \theta) \right)^2 dt \right)^{1/2}$$

$$\leq (R_2 - R_1)^{1/2} \left(\int_{R_1}^{R_2} \left(\int_{-\pi}^{\pi} \left(\frac{\partial v}{\partial r}(t, \theta) \right)^2 dt \right)^{1/2} \right).$$

Squaring both sides of (4.8) and integrating over Ω leads to the inequality

$$\begin{aligned} \iint_{\Omega} (v(r, \theta))^2 r dr d\theta &\leq (R_2 - R_1) \int_{-\pi}^{\pi} \int_{R_1}^{R_2} \left(\int_{-\pi}^{\pi} \left(\frac{\partial v}{\partial r}(t, \theta) \right)^2 dt \right) r dr d\theta \\ &= (R_2 - R_1) \frac{(R_2^2 - R_1^2)}{2} \int_{-\pi}^{\pi} \int_{R_1}^{R_2} \left(\frac{\partial v}{\partial r}(t, \theta) \right)^2 dt d\theta \\ &\leq \frac{(R_2 - R_1)(R_2^2 - R_1^2)}{2R_1} \iint_{\Omega} \left(\frac{\partial v}{\partial r}(t, \theta) \right)^2 t dt d\theta \end{aligned}$$

Hence,

$$\begin{aligned} \iint_{\Omega} v^2 dx dy &= \int_{-\pi}^{\pi} \int_{R_1}^{R_2} v^2 r dr d\theta \\ &\leq \frac{(R_2 - R_1)(R_2^2 - R_1^2)}{2R_1} \int_{-\pi}^{\pi} \int_{R_1}^{R_2} \left(\frac{\partial v}{\partial r} \right)^2 r dr d\theta \\ &\leq \frac{(R_2 - R_1)(R_2^2 - R_1^2)}{2R_1} \int_{-\pi}^{\pi} \int_{R_1}^{R_2} \left(\left(\frac{\partial v}{\partial r} \right)^2 + \frac{1}{r^2} \left(\frac{\partial v}{\partial \theta} \right)^2 \right) r dr d\theta \\ &= \frac{(R_2 - R_1)(R_2^2 - R_1^2)}{2R_1} \iint_{\Omega} \left(\left(\frac{\partial v}{\partial x} \right)^2 + \left(\frac{\partial v}{\partial y} \right)^2 \right) dx dy. \end{aligned}$$

Finally, the fact that $C_1^1(\bar{\Omega})$ is dense in $H_1^1(\Omega)$ guarantees that

(4.7) holds for all $v \in H_1^1(\Omega)$. This completes the proof of Lemma 4.4.

For all $v \in H_1^1(\Omega)$, the result (4.7) combined with the assumed positivity of g on $\partial_2\Omega$ shows that

$$\begin{aligned}
 (v, v)_1 &= \iint_{\Omega} (v^2 + (\frac{\partial v}{\partial x})^2 + (\frac{\partial v}{\partial y})^2) \, dx dy \\
 &\leq \left(\frac{(R_2 - R_1)(R_2^2 - R_1^2)}{2R_1} + 1 \right) \iint_{\Omega} ((\frac{\partial v}{\partial x})^2 + (\frac{\partial v}{\partial y})^2) \, dx dy \\
 &\leq \left(\frac{(R_2 - R_1)(R_2^2 - R_1^2)}{2R_1} + 1 \right) \left[\iint_{\Omega} ((\frac{\partial v}{\partial x})^2 + (\frac{\partial v}{\partial y})^2) \, dx dy + \int_{\partial_2\Omega} g v^2 \, d\sigma \right] \\
 &= \left(\frac{(R_2 - R_1)(R_2^2 - R_1^2)}{2R_1} + 1 \right) b(v, v).
 \end{aligned}$$

That is, $b(\cdot, \cdot)$ is $H_1^1(\Omega)$ -elliptic and the Lax-Milgram theorem now guarantees a unique solution of the variational Problem 4.2. The problem is now restate as: Find $T \in H_1^1(\Omega)$ such that

$$(4.9) \quad b(T, v) = \ell(v) \quad \text{for all } v \in H_1^1(\Omega).$$

Notice that $H_1^1(\Omega)$ is a very large space (it is dense in $L^2(\Omega)$) so that the variational problem (4.9) is, in general, impossible to solve analytically.

A viable alternative to obtaining the analytic solution of the problem (4.9) is the solution of the Ritz-Galerkin approximate

problem. The latter is formulated as follows: Find $T_n \in S_n$ such that

$$(4.10) \quad b(T_n, w) = \ell(w) \quad \text{for all } w \in S_n$$

where S_n is a finite dimensional subspace of $H_1^1(\Omega)$. The Lax-Milgram theorem again applies and guarantees a unique Ritz-Galerkin solution T_n . Suppose $\{\phi_i\}_{i=1}^n$ is a basis for the subspace S_n , then there exists scalars $\{a_i\}_{i=1}^n$ so that T_n can be expressed as

$$T_n = \sum_{i=1}^n a_i \phi_i.$$

Now, (4.10) can be rewritten as the linear system of equations

$$(4.11) \quad \sum_{i=1}^n a_i b(\phi_i, \phi_j) = \ell(\phi_j) \quad \text{for } j = 1, 2, \dots, n$$

for the determination of the coefficients $\{a_i\}_{i=1}^n$.

The Ritz-Galerkin solution T_n is an approximation to the solution T of the variational problem (4.9). To get a feeling for the accuracy of this approximation, notice that $b(\cdot, \cdot)$ defines a second inner product on $H_1^1(\Omega)$. Subtracting (4.10) from (4.9) shows that this inner product satisfies

$$b(T - T_n, w) = 0 \quad \text{for all } w \in S_n.$$

That is, T_n is the projection of T into the subspace S_n with respect

to the inner product $b(\cdot, \cdot)$. Recall that the projection satisfies

$$\|T - T_n\|_b = \inf\{\|T - w\|_b : w \in S_n\}$$

where $\|\cdot\|_b$ denotes the norm induced by $b(\cdot, \cdot)$. The continuity and ellipticity of $b(\cdot, \cdot)$ imply that $\|\cdot\|_1$ and $\|\cdot\|_b$ are equivalent norms, so there is a constant K , independent of the subspace S_n , such that

$$(4.12) \quad \|T - T_n\|_1 \leq K \inf\{\|T - w\|_1 : w \in S_n\}.$$

Hence, to obtain an accurate Ritz-Galerkin approximate, the subspace S_n is chosen so that $\inf\{\|T - w\|_1 : w \in S_n\}$ will be small.

4.3 Finite Element Approximation

In the context of the previous section, the finite element method is exactly the Ritz-Galerkin method equipped with special choices of the subspaces S_n . The construction of the most commonly used subspace S_n is briefly discussed in this section.

The annulus Ω is first approximated by a polygonal domain Ω_h depending on the positive number h (see Figure 4.6). Next, the region bounded by $\partial_2 \Omega_h$ is partitioned into triangles with maximum side length h (see Figure 4.7). A set of continuous functions $\{\phi_i\}$ is defined on the triangularized domain in such a way that each function ϕ_i is a

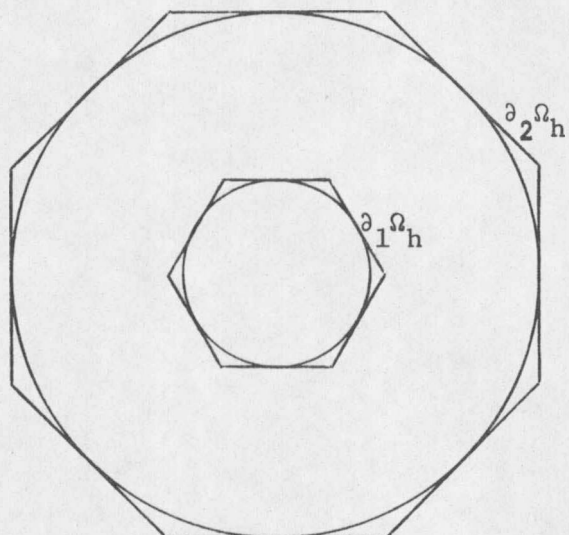


Figure 4.6. The approximation of the annulus Ω by the polygonal domain Ω_h .

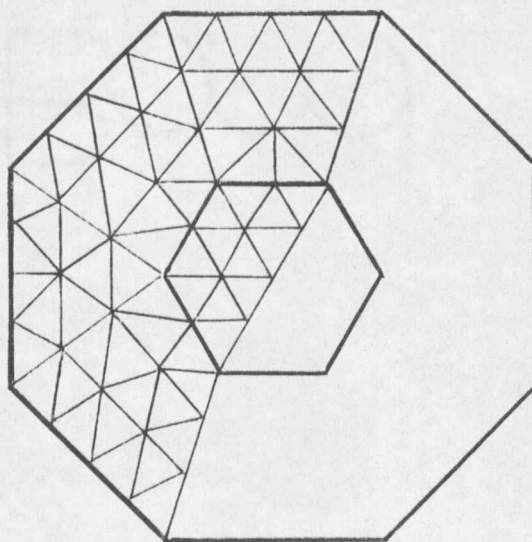


Figure 4.7. A partial triangularization of the domain bounded by $\partial_2 \Omega_h$.

piecewise polynomial which is zero over all but a small number of the triangular elements. The subspace S_n is taken as the span of a subset $\{\phi_i\}_{i=1}^n$ of these piecewise polynomials.

There are two difficulties with the finite element method as described above. Both difficulties occur near the boundaries of Ω and are results of the polygonal approximation of the curved boundaries.

To satisfy the essential boundary condition $T = 0$ on $\partial_1\Omega$, each basis function ϕ_i must vanish on $\partial_1\Omega_h$ and on the region between $\partial_1\Omega$ and $\partial_1\Omega_h$. Due to this, the infimum statement in (4.12) will in general not be small. There has been much work on this problem of essential boundary conditions on curved boundaries, and the survey paper of Bramble [7] discusses several approaches to this problem.

The second difficulty occurs on $\partial_2\Omega$. The discrete problem (4.11) contains integrals over the boundary $\partial_2\Omega$ and these integrals will in general be difficult to compute exactly. A numerical quadrature rule is usually employed in the finite element method to approximate these boundary integrals. However, the curved boundaries create difficulties which cannot be surmounted by the most accurate of quadrature rules. Due to the nature of the triangularized domain, there is the difficult problem of keeping track of the triangles the curve $\partial_2\Omega$ passes through. This is indeed a great difficulty, and many times to avoid it the boundary integrals are shifted to the polygonal

approximation $\partial_2 \Omega_h$ of $\partial_2 \Omega$. While this shifting of integrals may result in easier computations, it must be remembered that the boundary integrals implicitly involve the normal derivative of the trial functions ϕ_i . It is quite apparent that the normal derivatives with respect to the boundaries $\partial_2 \Omega_h$ and $\partial_2 \Omega$ do not often coincide. Indeed, the normal derivative with respect to the boundary $\partial_2 \Omega_h$ is undefined at the vertices of $\partial_2 \Omega_h$. Hence, due again to the difficulties caused by a curved boundary the infimum statement in (4.12) may not be small. There have been some techniques developed to address this problem of nonhomogeneous Neumann conditions on curved boundaries, and the reader is referred to the discussion of the isoparametric technique given by Strang and Fix [29] or the papers of Zlamal [34-36] on curved elements.

Although the techniques referred to can, in certain cases, give better approximations to solutions defined on domains with curved boundaries, they significantly increase the number of computations to be carried out to obtain the Ritz-Galerkin approximate T_n . What is needed is a technique that can handle the curved boundaries and remain efficient and easy to implement. The method of the next section appears to satisfy these requirements.

4.4 Annular Strip Approximation.

In this section, the annular strip method is proposed as an

alternative to the classical finite element method. This strip method avoids the boundary difficulties outlined in the previous section, yet the method retains certain salient features of the finite element method. Most notable among these features, from the point of view of numerical implementation, is the retention of a sparse discretized system for the efficient solution of (4.11).

As the basis functions for the method of the present section involve Fourier approximates, it will facilitate the discussion to have available the following two classical results on Fourier series [37].

Theorem 4.5. If f is continuous on the interval $[-\pi, \pi]$, then the Fourier series of f converges in mean to f on $[-\pi, \pi]$. That is

$$\int_{-\pi}^{\pi} (f(\theta) - (1/2 a_0 + \sum_{k=1}^N (a_k \cos k\theta + b_k \sin k\theta)))^2 d\theta$$

converges to zero as $N \rightarrow \infty$, where

$$a_k = \frac{1}{\pi} \int_{-\pi}^{\pi} f(\theta) \cos k\theta d\theta,$$

and

$$b_k = \frac{1}{\pi} \int_{-\pi}^{\pi} f(\theta) \sin k\theta d\theta.$$

Theorem 4.6. If f has m , $m \geq 2$, continuous derivatives on the

interval $[-\pi, \pi]$ and is 2π periodic, then the Fourier series of f converges uniformly to f in $[-\pi, \pi]$, and

$$|f(\theta) - (1/2 a_0 + \sum_{k=1}^N (a_k \cos k\theta + b_k \sin k\theta))| < K N^{-(m-1)}$$

for all $\theta \in [-\pi, \pi]$, where K is some positive constant.

Suppose $T(r, \theta)$ is a continuous function on the closed annulus

$$\bar{Q} = \{(r, \theta) : R_1 \leq r \leq R_2, -\pi < \theta \leq \pi\}$$

where (r, θ) are polar coordinates. For each r , $R_1 \leq r \leq R_2$, $T(r, \cdot)$ is continuous on the interval $[-\pi, \pi]$. This observation motivates the following approximation scheme.

First, define the set of nodes

$$R_1 = r_1 < r_2 < \cdots < r_n = R_2$$

in the radial direction by

$$r_i = R_1 + (i - 1) h,$$

where

$$h = (R_2 - R_1)/(n - 1).$$

These nodes partition the region $\bar{Q}$ into annular strips as displayed in Figure 4.8. For each r_i , approximate $T(r_i, \cdot)$ by the truncated Fourier

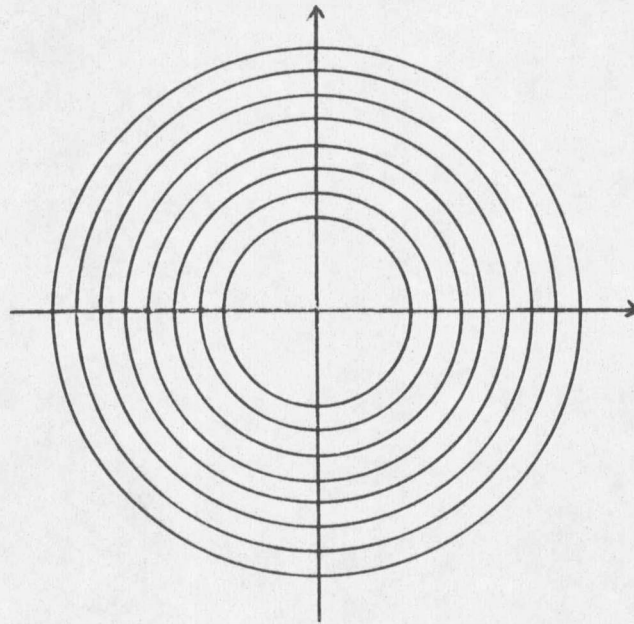


Figure 4.8. Ω partitioned into annular strips.

series

$$\frac{1}{2} a_{i,0} + \sum_{k=1}^N (a_{i,k} \cos k\theta + b_{i,k} \sin k\theta),$$

where

$$a_{i,k} = \frac{1}{\pi} \int_{-\pi}^{\pi} T(r_i, \theta) \cos k\theta \, d\theta$$

and

$$b_{i,k} = \frac{1}{\pi} \int_{-\pi}^{\pi} T(r_i, \theta) \sin k\theta \, d\theta.$$

These approximations are extended to the entire region $\bar{\Omega}$ by connecting the truncated Fourier series linearly across the annular strips. This is easily accomplished with the aid of the Chapeau functions $\{\Lambda_i\}_{i=1}^n$, which are geometrically defined via Figure 4.9. Thus, the approximation of the function T takes the form

$$T_{nN}(r, \theta) = \sum_{i=1}^n \left(\frac{1}{2} a_{i,0} + \sum_{k=1}^N (a_{i,k} \cos k\theta + b_{i,k} \sin k\theta) \right) \Lambda_i(r),$$

which will be termed the Chapeau-Fourier approximate of order (n, N) to the function T .

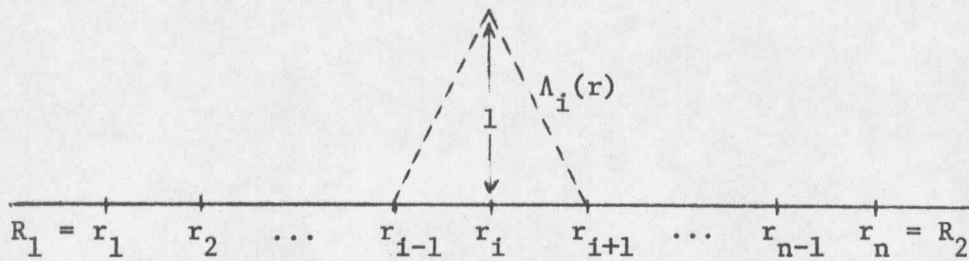


Figure 4.9. The Chapeau function Λ_i .

Theorem 4.5 is used to relate the Chapeau-Fourier approximate T_{nN} to a continuous function T on $\bar{\Omega}$ with respect to the Sobolev norm (2.3).

Theorem 4.7. Suppose the function T along with its first partials derivatives are continuous on the closed annulus $\bar{\Omega}$. Then there is a sequence $\{T_{nN}\}$ of Chapeau-Fourier approximates which converges to T in the Sobolev norm

$$||u||_1 = \left(\iint_{\Omega} (u^2 + (\frac{\partial u}{\partial x})^2 + (\frac{\partial u}{\partial y})^2) dx dy \right)^{1/2}.$$

Proof. Notice that by rewriting the Sobolev norm in polar coordinates

$$||u||_1 = \left(\int_{-\pi}^{\pi} \int_{R_1}^{R_2} (u^2 + (\frac{\partial u}{\partial r})^2 + \frac{1}{r^2} (\frac{\partial u}{\partial \theta})^2) r dr d\theta \right)^{1/2},$$

it need only be shown that given $\varepsilon > 0$, there is a Chapeau-Fourier approximate T_{nN} such that

$$||T - T_{nN}||_{\Omega} < \varepsilon, \quad ||\frac{\partial T}{\partial r} - \frac{\partial T_{nN}}{\partial r}||_{\Omega} < \varepsilon, \quad \text{and} \quad ||\frac{\partial T}{\partial \theta} - \frac{\partial T_{nN}}{\partial \theta}||_{\Omega} < \varepsilon$$

where $||\cdot||_{\Omega}$ denotes the L^2 -norm over Ω . Also notice that since $\bar{\Omega}$ is compact, T , $\frac{\partial T}{\partial r}$, and $\frac{\partial T}{\partial \theta}$ are all uniformly continuous on $\bar{\Omega}$.

To get the upper bound for $||T - T_{nN}||_{\Omega}$, the uniform continuity of T is used to find a $\delta_1 > 0$ such that

$$|T(r, \theta) - T(s, \phi)| < \varepsilon_1 = ((2\pi^{1/2} + 2^{1/2})(R_2^2 - R_1^2)^{1/2})^{-1} \varepsilon,$$

whenever $(r, \theta), (s, \phi) \in \Omega$ with $\text{dist}((r, \theta), (s, \phi)) < \delta_1$. Now, choose n_1 such that $h_1 = (R_2 - R_1)/(n_1 - 1) < \delta_1$ and define the nodes

$$r_i = R_1 + (i - 1) h_1 \quad \text{for } i = 1, \dots, n_1.$$

Then, for each i and each θ , $\text{dist}((r_i, \theta), (r_{i+1}, \theta)) < \delta_1$. Since the Fourier series of $T(r_i, \cdot)$ converges in mean to $T(r_i, \cdot)$ on $[-\pi, \pi]$ (Theorem 4.5), for all sufficiently large N_1 , the following inequality holds for each i

$$\|T(r_i, \theta) - (1/2 a_{i,0} + \sum_{k=1}^{N_1} (a_{i,k} \cos k\theta + b_{i,k} \sin k\theta))\|_I < \varepsilon_1,$$

where $\|\cdot\|_I$ denotes the L^2 -norm over the interval $I = [-\pi, \pi]$.

With n_1 and N_1 chosen as above, the triangle inequality yields

$$\|T - T_{n_1 N_1}\|_{\Omega} \leq \|T - \sum_{i=1}^{n_1} T(r_i, \cdot) \Lambda_i\|_{\Omega} + \|\sum_{i=1}^{n_1} T(r_i, \cdot) \Lambda_i - T_{n_1 N_1}\|_{\Omega}.$$

Upper bounds for the two right hand terms are calculated as follows:

$$\|T - \sum_{i=1}^{n_1} T(r_i, \cdot) \Lambda_i\|_{\Omega}^2 = \int_{R_1}^{R_2} \int_{-\pi}^{\pi} (T(r, \theta) - \sum_{i=1}^{n_1} T(r_i, \theta) \Lambda_i(r))^2 r d\theta dr$$

$$= \sum_{i=1}^{n_1-1} \int_{r_i}^{r_{i+1}} \int_{-\pi}^{\pi} (T(r, \theta) - T(r_i, \theta) \Lambda_i(r) - T(r_{i+1}, \theta) \Lambda_{i+1}(r))^2 r d\theta dr$$

$$= \sum_{i=1}^{n_1-1} \int_{r_i}^{r_{i+1}} \int_{-\pi}^{\pi} (T(r, \theta) - T(r_i, \theta) \left(\frac{r_{i+1} - r}{r_{i+1} - r_i} \right) - T(r_{i+1}, \theta) \left(\frac{r - r_i}{r_{i+1} - r_i} \right))^2 r d\theta dr$$

$$= \sum_{i=1}^{n_1-1} \int_{r_i}^{r_{i+1}} \int_{-\pi}^{\pi} \left(\left(\frac{r_{i+1} - r}{r_{i+1} - r_i} \right) (T(r, \theta) - T(r_i, \theta)) + \left(\frac{r - r_i}{r_{i+1} - r_i} \right) (T(r, \theta) - T(r_{i+1}, \theta)) \right)^2 r d\theta dr$$

$$= \sum_{i=1}^{n_1-1} \int_{r_i}^{r_{i+1}} \int_{-\pi}^{\pi} \left(\left(\frac{r_{i+1} - r}{r_{i+1} - r_i} \right)^2 (T(r, \theta) - T(r_i, \theta))^2 + 2 \left(\frac{r_{i+1} - r}{r_{i+1} - r_i} \right) \left(\frac{r - r_i}{r_{i+1} - r_i} \right) (T(r, \theta) - T(r_i, \theta)) (T(r, \theta) - T(r_{i+1}, \theta)) + \left(\frac{r - r_i}{r_{i+1} - r_i} \right)^2 (T(r, \theta) - T(r_{i+1}, \theta))^2 \right) r d\theta dr$$

$$< \sum_{i=1}^{n_1-1} \int_{r_i}^{r_{i+1}} \int_{-\pi}^{\pi} 4 \varepsilon_1^2 r d\theta dr = \int_{R_1}^{R_2} \int_{-\pi}^{\pi} 4 \varepsilon_1^2 r d\theta dr = 4\pi(R_2^2 - R_1^2) \varepsilon_1^2.$$

While,

$$\| \sum_{i=1}^{n_1} T(r_i, \cdot) \Lambda_i - T_{n_1 N_1} \|_{\Omega}^2 = \int_{R_1}^{R_2} \int_{-\pi}^{\pi} \left(\sum_{i=1}^{n_1} T(r_i, \theta) \Lambda_i(r) - T_{n_1 N_1}(r, \theta) \right)^2 r d\theta dr$$

$$= \sum_{i=1}^{n_1-1} \int_{r_i}^{r_{i+1}} \int_{-\pi}^{\pi} ((T(r_i, \theta) - (1/2 a_{i,0} + \sum_{k=1}^{N_1} (a_{i,k} \cos k\theta$$

$$+ b_{i,k} \sin k\theta))) \Lambda_i(r)$$

$$+ (T(r_{i+1}, \theta) - (1/2 a_{i+1,0} + \sum_{k=1}^{N_1} (a_{i+1,k} \cos k\theta$$

$$+ b_{i+1,k} \sin k\theta))) \Lambda_{i+1}(r))^2 r d\theta dr$$

$$\leq \sum_{i=1}^{n_1-1} \int_{r_i}^{r_{i+1}} \int_{-\pi}^{\pi} (|T(r_i, \theta) - (1/2 a_{i,0} + \sum_{k=1}^{N_1} (a_{i,k} \cos k\theta$$

$$+ b_{i,k} \sin k\theta))| |\Lambda_i(r)|$$

$$+ |T(r_{i+1}, \theta) - (1/2 a_{i+1,0} + \sum_{k=1}^{N_1} (a_{i+1,k} \cos k\theta$$

$$+ b_{i+1,k} \sin k\theta))| |\Lambda_{i+1}(r)|)^2 r d\theta dr$$

$$< \sum_{i=1}^{N_1} \int_{r_i}^{r_{i+1}} \int_{-\pi}^{\pi} (|T(r_i, \theta) - (1/2 a_{i,0} + \sum_{k=1}^{N_1} (a_{i,k} \cos k\theta + b_{i,k} \sin k\theta))|$$

$$+ |T(r_{i+1}, \theta) - (1/2 a_{i+1,0} + \sum_{k=1}^{N_1} (a_{i+1,k} \cos k\theta$$

$$+ b_{i+1,k} \sin k\theta))|)^2 r d\theta dr$$

$$\begin{aligned}
&= \sum_{i=1}^{n_1-1} \int_{r_i}^{r_{i+1}} || |T(r_i, \theta) - (1/2 a_{i,0} + \sum_{k=1}^{N_1} (a_{i,k} \cos k\theta + b_{i,k} \sin k\theta))| + \\
&\quad + |T(r_{i+1}, \theta) - (1/2 a_{i+1,0} + \sum_{k=1}^{N_1} (a_{i+1,k} \cos k\theta \\
&\quad + b_{i+1,k} \sin k\theta))| ||_I^2 r dr
\end{aligned}$$

$$\begin{aligned}
&\leq \sum_{i=1}^{n_1-1} \int_{r_i}^{r_{i+1}} (|| |T(r_i, \theta) - (1/2 a_{i,0} + \sum_{k=1}^{N_1} (a_{i,k} \cos k\theta + b_{i,k} \sin k\theta))| |_{\Gamma} \\
&\quad + || |T(r_{i+1}, \theta) - (1/2 a_{i+1,0} + \sum_{k=1}^{N_1} (a_{i+1,k} \cos k\theta \\
&\quad + b_{i+1,k} \sin k\theta))| |_{\Gamma})^2 r dr
\end{aligned}$$

$$\leq \sum_{i=1}^{n_1-1} \int_{r_i}^{r_{i+1}} 4 \varepsilon_1^2 r dr = 2(R_2^2 - R_1^2) \varepsilon_1^2.$$

Hence,

$$||T - T_{n_1 N_1}||_{\Omega} < 2(\pi(R_2^2 - R_1^2))^{1/2} \varepsilon_1 + (2(R_2^2 - R_1^2))^{1/2} \varepsilon_1 = \varepsilon.$$

By exactly the same argument as above, there is a Chapeau-Fourier approximate $(\frac{\partial T}{\partial \theta})_{n_2 N_2}$ of $\frac{\partial T}{\partial \theta}$ such that

$$||\frac{\partial T}{\partial \theta} - (\frac{\partial T}{\partial \theta})_{n_2 N_2}||_{\Omega} < \varepsilon.$$

Thus, it is sufficient to show that

$$\left(\frac{\partial T}{\partial \theta}\right)_{n_2 N_2} = \frac{\partial T_{n_2 N_2}}{\partial \theta}.$$

Notice that

$$\left(\frac{\partial T}{\partial \theta}\right)_{n_2 N_2} = \sum_{i=1}^{n_2} \left(\frac{1}{2} \tilde{a}_{i,0} + \sum_{k=1}^{N_2} (\tilde{a}_{i,k} \cos k\theta + \tilde{b}_{i,k} \sin k\theta) \right) \Lambda_i(r)$$

while

$$\frac{\partial T_{n_2 N_2}}{\partial \theta} = \sum_{i=1}^{n_2} \left(\sum_{k=1}^{N_2} (-ka_{i,k} \sin k\theta + kb_{i,k} \cos k\theta) \right) \Lambda_i(r).$$

By an elementary integration by parts, it can be shown that

$\tilde{a}_{i,k} = kb_{i,k}$ and $\tilde{b}_{i,k} = -ka_{i,k}$. For example, the first equality follows from

$$\begin{aligned} \tilde{a}_{i,k} &= \frac{1}{\pi} \int_{-\pi}^{\pi} \frac{\partial T}{\partial \theta}(r_i, \theta) \cos k\theta \, d\theta \\ &= \frac{1}{\pi} (T(r_i, \theta) \cos k\theta \Big|_{-\pi}^{\pi} + k \int_{-\pi}^{\pi} T(r_i, \theta) \sin k\theta \, d\theta) \\ &= k b_{i,k}. \end{aligned}$$

Similar calculations lead to the equality $\tilde{b}_{i,k} = -ka_{i,k}$.

To get the upper bound for $||\frac{\partial T}{\partial r} - \frac{\partial T}{\partial r} \frac{nN}{nN}||_{\Omega}$, the uniform continuity of $\frac{\partial T}{\partial r}$ is used to find a $\delta_3 > 0$ such that

$$|\frac{\partial T}{\partial r}(r, \theta) - \frac{\partial T}{\partial r}(s, \phi)| < \varepsilon_3 = ((\pi^{1/2} + \frac{2^{1/2}}{(R_2 - R_1)}) (R_2^2 - R_1^2)^{1/2})^{-1} \varepsilon,$$

whenever $\text{dist}((r, \theta), (s, \phi)) < \delta_3$. Now, choose n_3 such that

$h_3 = (R_2 - R_1)/(n_3 - 1) < \delta_3$ and define the nodes

$$r_i = R_1 + (i - 1) h_3 \quad \text{for } i = 1, \dots, n_3.$$

Again, since the Fourier series of $T(r_i, \cdot)$ converges in mean to $T(r_i, \cdot)$ on $[-\pi, \pi]$, for all sufficiently large N_3 the following inequality holds for each i .

$$||T(r_i, \theta) - (1/2 a_{i,0} + \sum_{k=1}^{N_3} (a_{i,k} \cos k\theta + b_{i,k} \sin k\theta))||_I < \frac{\varepsilon_3}{n_3 - 1}.$$

With n_3 and N_3 chosen as above, the triangle inequality yields

$$\begin{aligned} (4.13) \quad ||\frac{\partial T}{\partial r} - \frac{\partial T}{\partial r} \frac{n_3 N_3}{n_3 N_3}||_{\Omega} &\leq ||\frac{\partial T}{\partial r} - \sum_{i=1}^{n_3-1} h_3^{-1} (T(r_{i+1}, \cdot) - T(r_i, \cdot)) F_i||_{\Omega} \\ &+ ||\sum_{i=1}^{n_3-1} h_3^{-1} (T(r_{i+1}, \cdot) - T(r_i, \cdot)) F_i - \frac{\partial T}{\partial r} \frac{n_3 N_3}{n_3 N_3}||_{\Omega}, \end{aligned}$$

where the F_i are defined by

$$F_i(r) = \begin{cases} 1 & \text{if } r_i \leq r \leq r_{i+1} \\ 0 & \text{if } R_1 \leq r \leq r_i, r_{i+1} < r \leq R_2. \end{cases}$$

The Mean-Value Theorem is used to bound the first term on the right hand side of (4.13). For each $i = 1, 2, \dots, n_3-1$ and each θ , $-\pi < \theta \leq \pi$, there is a point $\xi_i(\theta)$ between r_i and r_{i+1} such that

$$\frac{\partial T}{\partial r}(\xi_i(\theta), \theta) = \frac{T(r_{i+1}, \theta) - T(r_i, \theta)}{r_{i+1} - r_i}.$$

Then

$$\begin{aligned} & \left\| \frac{\partial T}{\partial r} - \sum_{i=1}^{n_3-1} h^{-1}(T(r_{i+1}, \cdot) - T(r_i, \cdot)) F_i \right\|_{\Omega}^2 \\ &= \int_{R_1}^{R_2} \int_{-\pi}^{\pi} \left(\frac{\partial T}{\partial r}(r, \theta) - \sum_{i=1}^{n_3-1} h^{-1}(T(r_{i+1}, \theta) - T(r_i, \theta)) F_i(r) \right)^2 r d\theta dr \\ &= \sum_{i=1}^{n_3-1} \int_{r_i}^{r_{i+1}} \int_{-\pi}^{\pi} \left(\frac{\partial T}{\partial r}(r, \theta) - \frac{T(r_{i+1}, \theta) - T(r_i, \theta)}{r_{i+1} - r_i} \right)^2 r d\theta dr \\ &= \sum_{i=1}^{n_3-1} \int_{r_i}^{r_{i+1}} \int_{-\pi}^{\pi} \left(\frac{\partial T}{\partial r}(r, \theta) - \frac{\partial T}{\partial r}(\xi_i(\theta), \theta) \right)^2 r d\theta dr \end{aligned}$$

$$\langle \sum_{i=1}^{n_3-1} \int_{r_i}^{r_{i+1}} \int_{-\pi}^{\pi} \epsilon_3^2 r d\theta dr = \int_{R_1-\pi}^{R_2} \int_{-\pi}^{\pi} \epsilon_3^2 r d\theta dr = \pi(R_2^2 - R_1^2) \epsilon_3^2.$$

The second term on the right hand side of (4.13) is bounded as follows:

$$\begin{aligned} & \left\| \sum_{i=1}^{n_3-1} h_3^{-1} (T(r_{i+1}, \cdot) - T(r_i, \cdot)) F_i - \frac{\partial T_{n_3 N_3}}{\partial r} \right\|_{\Omega}^2 \\ &= \int_{R_1-\pi}^{R_2} \int_{-\pi}^{\pi} \left(\sum_{i=1}^{n_3-1} h_3^{-1} (T(r_{i+1}, \theta) - T(r_i, \theta)) F_i(r) - \frac{\partial T_{n_3 N_3}}{\partial r}(r, \theta) \right)^2 r d\theta dr \\ &= \sum_{i=1}^{n_3-1} \int_{r_i}^{r_{i+1}} \int_{-\pi}^{\pi} (h_3^{-1} (T(r_{i+1}, \theta) - T(r_i, \theta))) \\ & \quad - h_3^{-1} (1/2 a_{i+1,0} + \sum_{k=1}^{N_3} (a_{i+1,k} \cos k\theta + b_{i+1,k} \sin k\theta)) \\ & \quad - (1/2 a_{i,0} + \sum_{k=1}^{N_3} (a_{i,k} \cos k\theta + b_{i,k} \sin k\theta)))^2 r d\theta dr \\ &= \sum_{i=1}^{n_3-1} \int_{r_i}^{r_{i+1}} h_3^{-2} \left\| T(r_{i+1}, \theta) - (1/2 a_{i+1,0} + \sum_{k=1}^{N_3} (a_{i+1,k} \cos k\theta \right. \\ & \quad \left. + b_{i+1,k} \sin k\theta)) \right. \\ & \quad \left. - (T(r_i, \theta) - (1/2 a_{i,0} + \sum_{k=1}^{N_3} (a_{i,k} \cos k\theta + b_{i,k} \sin k\theta))) \right\|_I^2 r dr \end{aligned}$$

$$\begin{aligned}
& \leq \sum_{i=1}^{n_3-1} \int_{r_i}^{r_{i+1}} h_3^{-2} (||T(r_{i+1}, \theta) - (1/2 a_{i+1,0} + \sum_{k=1}^{N_3} (a_{i+1,k} \cos k\theta + b_{i+1,k} \sin k\theta))||_I \\
& \quad + ||T(r_i, \theta) - (1/2 a_{i,0} + \sum_{k=1}^{N_3} (a_{i,k} \cos k\theta + b_{i,k} \sin k\theta))||_I)^2 r dr \\
& < \sum_{i=1}^{n_3-1} \int_{r_i}^{r_{i+1}} h_3^{-2} \frac{4\epsilon_3^2}{(n_3 - 1)^2} r dr = \int_{R_1}^{R_2} \frac{4\epsilon_3^2}{R_1(R_2 - R_1)^2} r dr = \frac{2(R_2^2 - R_1^2)}{(R_2 - R_1)^2} \epsilon_3^2
\end{aligned}$$

Substituting the two upper bounds into (4.13) yields

$$||\frac{\partial T}{\partial r} - \frac{\partial T_{nN}}{\partial r}||_{\Omega} < (\pi(R_2^2 - R_1^2))^{1/2} \epsilon_3 + \frac{(2(R_2^2 - R_1^2))^{1/2}}{R_2 - R_1} \epsilon_3 = \epsilon.$$

Let $n = \max \{n_1, n_2, n_3\}$, it is clear that there exists an N large enough so that $||T - T_{nN}||_{\Omega} < \epsilon$, $||\frac{\partial T}{\partial r} - \frac{\partial T_{nN}}{\partial r}||_{\Omega} < \epsilon$, and $||\frac{\partial T}{\partial \theta} - \frac{\partial T_{nN}}{\partial \theta}||_{\Omega} < \epsilon$. This completes the proof of Theorem 4.7.

Theorem 4.6 can be used in an argument similar to the one above to prove the following theorem.

Theorem 4.8. If $T \in C^m(\bar{\Omega})$, $m \geq 3$, then

$$||T - T_{nN}||_1 = O(N^{-(m-2)}) + O(h)$$

where $h = \frac{R_2 - R_1}{n - 1}$ is the mesh size of the partition in the radial direction.

Returning now to the variational problem, it is seen that the annular strip method, like the finite element method, is exactly the Ritz-Galerkin method with a special choice of the subspaces S_n . The Ritz-Galerkin approximates are taken from the spaces

$$S_{nN} = \text{span}\{\Lambda_i(r), \Lambda_i(r)\cos k\theta, \Lambda_i(r)\sin k\theta\}$$

where $i = 2, \dots, n$ and $k = 1, \dots, N$. Let $\hat{T}_{nN}$ denote the solution from S_{nN} of the Ritz-Galerkin problem (4.10) and recall (Equation 4.12) that T_{nN} satisfies

$$(4.14) \quad \|T - \hat{T}_{nN}\|_1 \leq K \inf\{\|T - w\|_1 : w \in S_{nN}\}$$

where T is the true solution of problem (4.9) and K is a positive constant independent of S_{nN} . Now, since the Chapeau-Fourier approximate T_{nN} is an element of S_{nN} , Theorem 4.7 or Theorem 4.8 can be used in combination with (4.14) to guarantee the convergence of the annular strip method.

The use of annular strips instead of triangles (the typical finite elements), avoids the boundary difficulties which were encountered in Section 4.3. Specifically, every function in S_{nN} is zero on $\partial_1\Omega$ but not on a whole region near $\partial_1\Omega$. Also, the boundary integrals in the discrete problem (4.11) are essentially the Fourier coefficients of the function g defined on $\partial_2\Omega$. Therefore, if these

Fourier coefficients are known or have been accurately approximated beforehand, the algebraic system (4.11) which defines the Ritz-Galerkin approximate is very easy to assemble.

For the problem (4.1) - (4.3), there is one further simplifying hypothesis which can be made about the function g . In this direction, parametrize the curve $\partial_2\Omega$ by

$$x = R_2 \cos \theta \text{ and } y = R_2 \sin \theta \quad (-\pi < \theta \leq \pi).$$

Due to physical considerations, it can be assumed that g is an even function defined on the interval $[-\pi, \pi]$. Hence, the solution $T(r, \theta)$ will be an even function on the interval $[-\pi, \pi]$ for each fixed r , $R_1 \leq r \leq R_2$. The Ritz-Galerkin approximate can therefore be selected from the space

$$S_{nN} = \text{span}\{\Lambda_i(r) \cos j\theta\},$$

where $i = 2, \dots, n$ and $j = 0, \dots, N$. Let the basis elements of S_{nN} be denoted by

$$\phi_{ij}(r, \theta) = \Lambda_i(r) \cos j\theta,$$

then the Ritz-Galerkin approximate takes the form

$$\hat{T}_{nN} = \sum_{i=2}^n \sum_{j=0}^N a_{ij} \phi_{ij}.$$

With this notation, the discrete problem (4.11) is expressed as

$$(4.15) \quad \sum_{i=2}^n a_{ij} b(\phi_{ij}, \phi_{km}) = \ell(\phi_{km})$$

for $k = 2, \dots, n$ and $m = 0, \dots, N$.

Now, suppose the function g has the Fourier expansion

$$g(\theta) \approx 1/2 g_0 + \sum_{k=1}^{\infty} g_k \cos k\theta.$$

Then the linear algebraic system (4.15) is assembled as follows:

$$\ell(\phi_{km}) = \int_{\partial\Omega_2} (k_1 g + k_2) \phi_{km} d\sigma = R_2 \Lambda_k(R_2) \int_{-\pi}^{\pi} (k_1 g(\theta) + k_2) \cos m\theta d\theta,$$

so that

$$\ell(\phi_{km}) = \begin{cases} 2\pi R_2 (k_1 g_0 + k_2) & k = n, m = 0 \\ \pi R_2 k_1 g_m & k = n, m = 1, \dots, N \\ 0 & \text{otherwise.} \end{cases}$$

$$b(\phi_{ij}, \phi_{km}) = \iint_{\Omega} \left(\frac{\partial \phi_{ij}}{\partial x} \frac{\partial \phi_{km}}{\partial x} + \frac{\partial \phi_{ij}}{\partial y} \frac{\partial \phi_{km}}{\partial y} \right) dx dy + \int_{\partial_2 \Omega} g \phi_{ij} \phi_{km} d\sigma$$

$$= \int_{R_1}^{R_2} \int_{-\pi}^{\pi} \left(\frac{\partial \phi_{ij}}{\partial r} \frac{\partial \phi_{km}}{\partial r} + \frac{1}{r^2} \frac{\partial \phi_{ij}}{\partial \theta} \frac{\partial \phi_{km}}{\partial \theta} \right) r d\theta dr$$

$$+ 1/2 R_2 \Lambda_i(R_2) \Lambda_k(R_2) \int_{-\pi}^{\pi} (g(\theta) \cos(j+m)\theta + g(\theta) \cos(j-m)\theta) d\theta,$$

so that with $h = \frac{R_2 - R_1}{n - 1}$ and $r_i = R_1 + ih$,

$$b(\phi_{ij}, \phi_{km}) = \begin{cases} \frac{\pi}{h} (r_n + r_{n-1}) + 2\pi R_2 g_0 & i = k = n, j = m = 0 \\ \left[\frac{\pi}{2h} (r_n + r_{n-1}) + \frac{\pi m^2}{2h} \left[\frac{1}{2} r_n^2 - 2r_{n-1}r_n + \frac{3}{2} r_{n-1}^2 \right] \right. & i = k = n \\ \quad \left. + r_{n-1} \log(r_n/r_{n-1}) \right] + \frac{1}{2} \pi R_2 [g_{2m} + 2g_0] & j = m \neq 0 \\ \frac{1}{2} \pi R_2 [g_{j+m} + g_{|j-m|}] & i = k = n, j \neq m \\ \frac{4\pi}{h} r_k & i = k \neq n, j = m = 0 \\ \left[\frac{2\pi}{h} r_k + \frac{\pi m^2}{h^2} \left[\frac{3}{2} (r_{k-1}^2 - r_{k+1}^2) + 4hr_k \right] \right. & i = k \neq n \\ \quad \left. + r_{k+1}^2 \log(r_{k+1}/r_k) + r_{k-1}^2 \log(r_k/r_{k-1}) \right] & j = m \neq 0 \\ -\frac{\pi}{h} (r_k + r_{k-1}) & i = k - 1, j = m = 0 \\ \left[-\frac{\pi}{2h} (r_k + r_{k-1}) + \frac{\pi m^2}{h^2} \left[\frac{1}{2} (r_k^2 - r_{k-1}^2) \right] \right. & i = k - 1 \\ \quad \left. + r_{k-1} r_k \log(r_{k-1}/r_k) \right] & j = m \neq 0 \\ -\frac{\pi}{h} (r_{k+1} + r_k) & i = k + 1, j = m = 0 \\ \left[-\frac{\pi}{2h} (r_{k+1} + r_k) + \frac{\pi m^2}{h^2} \left[\frac{1}{2} (r_{k+1}^2 - r_k^2) \right] \right. & i = k + 1 \\ \quad \left. + r_k r_{k+1} \log(r_k/r_{k+1}) \right] & j = m \neq 0 \\ 0 & \text{otherwise.} \end{cases}$$

It is apparent that the algebraic system (4.15) is very easy to assemble and, due to the orthogonality of the Chapeau functions, this system has a sparse symmetric coefficient matrix. In Example 4.9 of the next section, it is shown how to order the coefficients a_{ij} of (4.14) so that the resulting system can be efficiently solved.

4.5 Numerical Implementation and Examples

In this section, the annular strip method is applied to two test examples. The first problem is of the form given by (4.1) - (4.3). This example is used to exhibit the implementation and accuracy of the annular strip method. The second example is taken from a paper by Bruch and Sloss [8]. This second example is included to demonstrate the application of the strip method to a Poisson problem.

The Bruch-Sloss paper addresses the problem of approximating the solution of Laplace's equation with essential boundary conditions on certain doubly connected domains (in particular, an annulus). The Bruch-Sloss procedure is to reformulate the problem as an integral equation on the outer boundary of the doubly connected domain, and then to use a finite element scheme to solve this integral equation. The Bruch-Sloss procedure, like the annular strip method, avoids the boundary difficulties of the classical finite element scheme of

Section 4.3

To solve the Bruch-Sloss problem via the annular strip method, the problem is reformulated as a Poisson problem with homogeneous boundary conditions. The annular strip method is seen to provide very accurate approximations at a low cost. In fact, an annular strip approximation which is comparable to the results of Bruch and Sloss was obtained at a fraction of the computer cost reported for the Bruch-Sloss procedure. This is consistent with the opinion [9,10,30] that the finite strip method provides a cheap (with regards to computer costs) means of obtaining accurate approximations to solutions of problems defined on regular geometric domains.

Example 4.9. As a first example, consider the problem of finding the function $T(r,\theta)$ which satisfies

$$\Delta T = 0 \quad \text{on } \Omega,$$

$$T = 0 \quad \text{on } \partial_1 \Omega,$$

and

$$\frac{dT}{dn} + \frac{5(2 - \cos \theta)}{6(2 + \cos \theta)} T = (3 + 5 \log 2) \frac{5(2 - \cos \theta)}{6(2 + \cos \theta)} \quad \text{on } \partial_2 \Omega,$$

where the Ω and $\partial_1 \Omega$ are defined as in Figure 4.5 with $R_1 = 1$ and $R_2 = 2$. This problem has the known solution

$$T(r, \theta) = 5 \log r + (r^{-1} + r) \cos \theta$$

The approximate Fourier expansion

$$\begin{aligned} g(\theta) \approx & 1.0912 - 1.0313 \cos \theta + .2763 \cos 2\theta - .074 \cos 3\theta \\ & + .0198 \cos 4\theta - .0053 \cos 5\theta + .0014 \cos 6\theta \end{aligned}$$

of the function

$$g(\theta) = \frac{5(2 - \cos \theta)}{6(2 + \cos \theta)}$$

is used in the assembly of the algebraic system (4.15) (described at the end of the previous section), which determines an annular strip approximation

$$\hat{T}_{nN} = \sum_{i=2}^n \sum_{j=0}^N a_{ij} \phi_{ij}$$

to the true solution T . Figure 4.10 is an example of the algebraic system (4.15) for the case $n = 5$ and $N = 2$. All numbers are rounded off to the first decimal place to facilitate the display in Figure 4.10. Notice the symmetry of the system and that by arranging the unknown coefficients a_{ij} in ascending order according to the index i first and the index j second, the first 9 columns of the matrix in Figure 4.10 have at most 3 nonzero entries and that these entries occur at the (k, k) and $(k \pm 3, k)$ positions. Due to the lower right

$$\begin{bmatrix}
 20.0 & 0 & 0 & -11.0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\
 0 & 10.1 & 0 & 0 & -5.5 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\
 0 & 0 & 10.5 & 0 & 0 & -5.4 & 0 & 0 & 0 & 0 & 0 & 0 \\
 -11.0 & 0 & 0 & 24.0 & 0 & 0 & -13.0 & 0 & 0 & 0 & 0 & 0 \\
 0 & -5.5 & 0 & 0 & 12.1 & 0 & 0 & -6.5 & 0 & 0 & 0 & 0 \\
 0 & 0 & -5.4 & 0 & 0 & 12.4 & 0 & 0 & -6.4 & 0 & 0 & 0 \\
 0 & 0 & 0 & -13.0 & 0 & 0 & 28.0 & 0 & 0 & -15.0 & 0 & 0 \\
 0 & 0 & 0 & 0 & -6.5 & 0 & 0 & 14.1 & 0 & 0 & -7.5 & 0 \\
 0 & 0 & 0 & 0 & 0 & -6.4 & 0 & 0 & 14.4 & 0 & 0 & -7.4 \\
 0 & 0 & 0 & 0 & 0 & 0 & -15.0 & 0 & 0 & 19.4 & -2.1 & .6 \\
 0 & 0 & 0 & 0 & 0 & 0 & 0 & -7.5 & 0 & -2.1 & 10.0 & -1.1 \\
 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & -7.4 & .6 & -1.1 & 9.9
 \end{bmatrix}
 \begin{bmatrix}
 a_{20} \\
 a_{21} \\
 a_{22} \\
 a_{30} \\
 a_{31} \\
 a_{32} \\
 a_{40} \\
 a_{41} \\
 a_{42} \\
 a_{50} \\
 a_{51} \\
 a_{52}
 \end{bmatrix}
 =
 \begin{bmatrix}
 0 \\
 0 \\
 0 \\
 0 \\
 0 \\
 0 \\
 0 \\
 0 \\
 0 \\
 28.2 \\
 -13.3 \\
 3.6
 \end{bmatrix}$$

Figure 4.10. An example of the algebraic system for the determination of an annular strip approximation.

hand submatrix, the last 3 columns of the matrix in Figure 4.10 do not follow the pattern of the previous columns. This submatrix is a result of the boundary integrals in (4.15). In general, if m Cheapeau functions and M trigonometric functions are used in the annular strip procedure, then the algebraic system (4.15), with the unknown coefficients α_{ij} arranged as above, will have a symmetric $(mM) \times (mM)$ coefficient matrix of which the first $(m-1)M$ columns will have nonzero entries occurring only in the (k,k) and $(k \pm M,k)$ positions. It is clear that a Gaussian elimination procedure, written to take advantage of the special form of this sparse coefficient matrix, can efficiently solve the algebraic system (4.15). In fact, instead of the usual $O((mM)^3)$ operations in the general Gaussian elimination procedure, matters have been arranged so that the operation count is on the order $O((m-1)M + M^3)$.

Annular strip approximations $\hat{T}_{nN}$ to the solution T of the present problem are computed with various choices for n and N and bounds for the uniform error of these approximations are listed in Table 4.1. These bounds are calculated as follows: Recall that the annular strip approximate $\hat{T}_{nN}$ is defined on the partitioned domain in Figure 4.8. Also recall that the annular strips of Figure 4.8 are defined by the subdivision

Table 4.1. Annular strip approximation to the solution of Example 4.9:

$$\Delta T = 0 \quad \text{on } \Omega,$$

$$T = 0 \quad \text{on } \partial_1 \Omega,$$

$$\frac{dT}{dn} + \frac{5(2 - \cos \theta)}{6(2 + \cos \theta)} T = (3 + 5 \log 2) \frac{5(2 - \cos \theta)}{6(2 + \cos \theta)} \quad \text{on } \partial_2 \Omega.$$

N	n	h	Uniform Error Bound
3	3	.5	.240774
2	5	.25	.060252
2	9	.125	.014947
3	17	.0625	.0038161

$$1 = r_1 < r_2 < \dots < r_n = 2$$

of the radial interval $[1,2]$ where

$$r_i = 1 + (i - 1)h \quad \text{with } h = 1/(n - 1).$$

An application of the triangle inequality gives a bound on the uniform error of the annular strip approximate $\hat{T}_{nN}$,

$$(4.16) \quad ||T - \hat{T}_{nN}||_{\infty} \leq ||T - \sum_{i=2}^n T(r_i, \cdot) \Lambda_i||_{\infty} + ||\sum_{i=2}^n T(r_i, \cdot) \Lambda_i - \hat{T}_{nN}||_{\infty}$$

where Λ_i is the Cheapeau function defined in the radial direction based at the node r_i . Thus, the bound on the uniform error is composed of two parts. The first part $||T - \sum_{i=2}^n T(r_i, \cdot) \Lambda_i||_{\infty}$ is a measure of the linear approximation to T in the radial direction, and the second part $||\sum_{i=2}^n T(r_i, \cdot) \Lambda_i - \hat{T}_{nN}||_{\infty}$ is a measure of the trigonometric polynomial approximation to $T(r_i, \cdot)$ for each $i = 2, 3, \dots, n$.

It should be pointed out that in all cases the uniform error bound is dominated by the linear approximation error term of (4.16). In fact, in all cases the linear approximation error comprises at least 90% of the uniform error bounds reported in Table 4.1. This explains the apparent $O(h^2)$ convergence of the error bounds. Such convergence is expected for linear approximations of functions with two continuous derivatives [29].

As indicated above, the second term of the uniform error bound (4.16) was very small. That is, in all cases the trigonometric polynomial approximation $\hat{T}_{nN}(r_i, \cdot)$ to $T(r_i, \cdot)$ is very accurate. As an example, consider the case $n = 9$ and $N = 2$. The true solution T at the radial node $r_5 = 1.5$ is the trigonometric polynomial

$$T(1.5, \theta) = 2.027326 - .833333 \cos \theta,$$

while the annular strip approximate $\hat{T}_{92}$ at the radial node $r_5 = 1.5$ is the trigonometric polynomial

$$\hat{T}_{92}(1.5, \theta) = 2.026381 - .833378 \cos \theta + .000012 \cos 2 \theta.$$

Such accuracy is seen to be typical for the annular strip method applied to the present problem.

The observations of the previous two paragraphs suggest that more accurate approximations can be obtained by using smoother approximations (cubic splines for example) in the radial direction. However, the use of smoother radial functions results in a more complex algebraic system (4.11). Therefore, due to the simplicity and efficient solvability of the system (4.15), this author recommends that the Cheapeau function be retained in the annular strip method and that convergence be induced by decreasing the step size h in the radial direction.

Example 4.10. As a second example, consider the problem of finding the function $u(r, \theta)$ which satisfies

$$\Delta u = 0 \quad \text{on } \Omega,$$

$$u = \cos 2\theta \quad \text{on } \partial_1 \Omega,$$

and

$$u = 1 \quad \text{on } \partial_2 \Omega$$

where the Ω and $\partial_1 \Omega$ are defined as in Figure 4.5 with $R_1 = 1$ and $R_2 = 2$. In order to apply the annular strip method, the present problem is reformulated as a Poisson problem with homogeneous boundary conditions. In this direction set

$$T(r, \theta) = u(r, \theta) - (r - 1) + (r - 2) \cos 2\theta.$$

The problem can now be restated as: Find $T(r, \theta)$ which satisfies

$$(4.17) \quad -\Delta T = f(r, \theta) = + \left(\frac{3}{r} + \frac{8}{r^2} \right) \cos 2\theta \quad \text{on } \Omega$$

and

$$(4.18) \quad T = 0 \quad \text{on } \partial_1 \Omega \cup \partial_2 \Omega.$$

This problem has the known solution

$$T(r, \theta) = \frac{\log r}{\log 2} - (r - 1) - \left(\frac{r^2}{15} - r + 2 - \frac{16}{15r^2} \right) \cos 2\theta.$$

The equivalent variational form of (4.17) and (4.18) is [16,19]: Find $T \in H_0^1(\Omega)$ such that

$$b(T, v) = \ell(v) \quad \text{for all } v \in H_0^1(\Omega)$$

where the bilinear form $b(\cdot, \cdot)$ is defined by

$$b(T, v) = \iint_{\Omega} \left(\frac{\partial T}{\partial x} \frac{\partial v}{\partial x} + \frac{\partial T}{\partial y} \frac{\partial v}{\partial y} \right) dx dy \quad \text{for } T, v \in H_0^1(\Omega)$$

and the linear functional ℓ is defined by

$$\ell(v) = \iint_{\Omega} f v \, dx dy \quad \text{for } v \in H_0^1(\Omega).$$

Using the notation of the previous section, an annular strip approximation $\hat{T}_{nN}$ to the true solution T is obtained by solving the Ritz-Galerkin problem: Find $\hat{T}_{nN} \in S_{nN}$ such that

$$(4.19) \quad b(\hat{T}_{nN}, v) = \ell(v) \quad \text{for all } v \in S_{nN}.$$

The requirement (4.19) produces an algebraic system for the determination of $\hat{T}_{nN}$ which is similar in form to that of the previous example. Bounds for the uniform error of the Ritz-Galerkin approximations $\hat{T}_{nN}$ are calculated and listed in Table 4.2. Once again, the linear approximation in the radial direction is responsible

Table 4.2. Annular strip approximation to the solution of Example 4.10:

$$-\Delta T = \frac{1}{r} + \left(\frac{3}{r} - \frac{1}{r^2}\right) \cos 2 \quad \text{on } \Omega$$

$$T = 0 \quad \text{on } \partial_1 \Omega \cup \partial_2 \Omega.$$

N	n	h	Uniform Error Bound
$N \geq 2$	3	.5	.2924724
$N \geq 2$	5	.25	.0730168
$N \geq 2$	9	.125	.0182703
$N \geq 2$	17	.0625	.0046034

for the major portion (over 98%) of the error bounds listed in Table 4.2. This implies that the trigonometric polynomial approximations to each $T(r_i, \cdot)$ must be very accurate. In fact, for the case $n = 9$ and $N \geq 2$, the true solution T at the radial node $r_5 = 1.5$ is the trigonometric polynomial

$$T(r_5, \theta) = .084963 - .175926 \cos 2\theta,$$

while the annular strip approximate $\hat{T}_{9N}$ at r_5 is the trigonometric polynomial

$$\hat{T}_{9N}(r_5, \theta) = .084853 - .175909 \cos 2\theta.$$

Such accuracy is again typical for the annular strip method as it applies to the present problem.

5. SUMMARY

The investigations of this thesis have addressed the numerical solution of two elliptic boundary value problems which arise in the mathematical model of two well known physical phenomena. In each instance, it has been shown that the pertinent physical quantity being approximated can be feasibly obtained (accurately and economically) with the proposed methods. These final remarks summarize the procedures and specify extensions that the author feels warrant further examination.

The radial Schrodinger eigenvalue problem was defined and studied in Chapter 3. This Sturm-Liouville problem is classified as a singular problem for two distinct reasons. Firstly, there is the infinite domain of definition ($\mathbb{R}_+$) over which the solution is sought and, due to the fact that the equation admits zero as a regular singular point, the eigenfunctions can be, and in instances are, singular at the origin. It is well known that these types of singularities pose difficulties for standard numerical procedures such as the finite difference and finite element methods. It has been shown that the sinc-collocation method of Chapter 3 surmounts the above difficulties for the radial Schrodinger problem. Indeed, the methods of Chapter 3 are applicable to Sturm-Liouville problems defined on infinite intervals in general. With regard to singular

eigenfunctions, the error bound obtained in Theorem 3.2 is valid whether or not the eigenfunctions are singular at zero (recall, in particular, Example 3.5).

It was further noted in the examples of Chapter 3 that the smaller eigenvalues of the sinc-collocation discrete problem appear to provide good approximations to the lower end of the spectrum of the continuous problem. A result supporting these numerical observations is of great interest to this author. In particular, it would be interesting to have an example which analytically corroborates the numerical evidence of Chapter 3. This type of example might then shed light on how a general result could be obtained. Indeed, the Hermite equation of Example 3.6 was included solely to lend credence to these remarks and the following conjecture.

Conjecture 5.1. Let $q(x)$ be nonnegative for each $x \in \mathbb{R}$. Denote by

$$\lambda_1 \leq \lambda_2 \leq \cdots \leq \lambda_{2N+1}$$

the first $2N + 1$ eigenvalues of the following eigenvalue problem:

$$-u''(x) + q(x)u(x) = \lambda u(x) \quad \text{for } x \in \mathbb{R}$$

subject to the boundary conditions

$$\lim_{x \rightarrow \pm\infty} u(x) = 0.$$

Then the $(2N + 1) \times (2N + 1)$ system

$$\left[\frac{1}{h^2} I_N^{(2)} + D_N(q(x)) \right] \vec{v}_N = \lambda \vec{v}_N$$

(defined in Chapter 3) has $2N + 1$ distinct eigenvalues $\{\Lambda_i\}_{i=1}^{2N+1}$ such that Λ_i approximates λ_i for all $i = 1, 2, \dots, 2N + 1$.

Investigation by this author of the discretized Hermite equation has already led to several interesting linear algebra problems. In particular, the theory of Toeplitz matrices and the circle of ideas centering around the notions of symmetric-centrosymmetric matrices have quite naturally arisen, but the present theory here does not seem to answer the conjecture. This author is confident that further studies in these areas of linear algebra will eventually lead to the resolution of the conjecture for the Hermite equation and, hopefully, to the general Schrodinger equation.

As stated earlier, the sinc function techniques of Section 2.3 have been the basis for many numerical schemes. However, the sinc-collocation method of Chapter 3 is the first application of these techniques to the numerically complex problem of approximating the eigenvalues of differential or integral equations. This large group of eigenvalue problems provide other possibilities for the application of the sinc function techniques, and in particular, the sinc-

collocation method.

The second elliptic boundary value problem of this thesis is a heat transfer problem defined on an annular domain. The classical solution of this problem is a harmonic function, and therefore, under different circumstances (a different domain) would lend itself well to approximation by piecewise polynomials such as those obtained by the finite element method. However, the curved boundaries of the domain of the problem pose difficulties for the finite element method and the accuracy of the resulting piecewise polynomial approximation will in general not be satisfactory relative to the work required for the acquisition of the approximation. The annular strip method of Chapter 4 is a viable alternative to the finite element method in that the former retains the salient features of the finite element method and avoids the boundary difficulties via a judicious choice of trial functions.

A distinctive feature of the annular strip method is the simplicity in the construction and subsequent solution of the discrete problem from which the annular strip approximation is obtained. This simplicity is easily traced to the orthogonality of the trigonometric functions and the semi-orthogonality of the Cheapeau functions.

As was noted in the examples of Chapter 4, the major portion of the error in the annular strip approximations is due to the linear

approximation (use of Cheapeau functions) in the radial direction. An alternative to the use of Cheapeau functions in the radial direction is the use of smoother radial functions such as cubic B-splines. It is clear that this latter choice of radial functions would produce a more complex discrete problem and, like the finite element procedure, the accuracy of the resulting approximation might not be satisfactory relative to the work required for the acquisition of the approximation. A thorough analysis of accuracy in approximation relative to computer cost for the annular strip procedure with both linear and cubic spline radial functions is an interesting direction for further investigations.

As stated earlier, the ideas behind the annular strip method are not new. Indeed, Fourier expansions coupled with spline approximations have been used by the engineering community in a procedure referred to as the finite strip method. However, the convergence results obtained for the annular strip method in Chapter 4 appear to be the first such results for any finite strip-like procedure. This author believes that the convergence results for the annular strip method can be extended to other applications of the finite strip procedure. In fact, the convergence results of Chapter 4 apply, with only a minor modification, to establish convergence of the examples reported in Chakrabarti's investigation of heat conduction in

plates by the finite strip method [9].

In the direction of the extension of the annular strip method to higher dimensions, this author feels that the application of the notions and theorems of Chapter 4 to the problem (4.1) - (4.3) in $\mathbb{R}^n$, and in particular $\mathbb{R}^3$ (the domain now being concentric balls), warrants attention for both analytic and numerical reasons. Once again, the variational formulation of this problem is nested in the theory of Sobolev spaces as presented in Chapter 2. In the $\mathbb{R}^3$ situation, Laplace's equation $\Delta T = 0$ (Equations (4.1)) is expressed in spherical coordinates as

$$r^2 \sin \theta \frac{\partial^2 T}{\partial r^2} + 2r \sin \theta \frac{\partial T}{\partial r} + \sin \theta \frac{\partial^2 T}{\partial \theta^2} + \cos \theta \frac{\partial T}{\partial \theta} + \frac{1}{\sin \theta} \frac{\partial^2 T}{\partial \phi^2} = 0.$$

Any solution $T(r, \theta, \phi)$ of this equation is known as a spherical harmonic. With regards to the approximation of a spherical harmonic, instead of the obvious product Fourier series angular approximation as in Chapter 4, it is well known that the angular solution admits an expression in terms of zonal harmonics, and in particular, representation in terms of the Legendre polynomials. Zonal harmonic angular approximations combined with linear radial approximations should, as in the two dimensional problems, provide an accurate yet efficient means of approximation.

BIBLIOGRAPHY

1. Adams, R. A., Sobolev Spaces, Academic Press, New York, 1975.
2. Agmon, S., Elliptic Boundary Value Problems, Van Nostrand Company, Inc., New York, 1965.
3. Aubin, J. P., Approximation of Elliptic Boundary-Value Problems, John Wiley and Sons, Inc., New York, 1972.
4. Babuska, I. and A. K. Aziz, Survey Lectures on the Mathematical Foundations of the Finite Element Method, in Aziz, A. K., Ed., The Mathematical Foundations of the Finite Element Method with Applications to Partial Differential Equations, Academic Press, New York, 1972, pp. 5-359.
5. Birkhoff, G., C. DeBoor, B. Swartz, and B. Wendroff, Rayleigh-Ritz Approximation by Piecewise Cubic Polynomials, SIAM J. Numer. Anal., 3 (1966), pp. 188-203.
6. Birkhoff, G. and G. C. Rota, Ordinary Differential Equations, third edition, John Wiley and Sons, Inc., New York, 1978.
7. Bramble, J. H., A Survey of Some Finite Element Methods Proposed for Treating the Dirichlet Problem, Advances in Math., 16 (1975), pp. 187-196.
8. Bruch, J. C. and J. M. Sloss, Harmonic Approximation with Dirichlet Data on Doubly Connected Regions, SIAM J. Numer. Anal., 14 (1977), pp. 994-1005.
9. Chakrabarti, S., Heat Conduction in Plates by Finite Strip Method, J. Engrg. Mech. Div., ASCE 106 EM2 (1980), pp. 233-244.
10. Cheung, Y. K., Finite Strip Method in Structural Analysis, Pergamon Press, New York, 1976.
11. Ciarlet, P. G., M. H. Schultz, and R. S. Varga, Numerical Methods of High-Order Accuracy for Nonlinear Boundary Value Problems: II. Eigenvalue Problems, Numer. Math., 12 (1968), pp. 120-133.

12. Johnson, O. G., Error Bounds for Sturm-Liouville Eigenvalue Approximations by Several Piecewise Cubic Rayleigh-Ritz Methods, SIAM J. Numer. Anal., 6 (1969), pp. 317-333.
13. Lund, J. R., Numerical Evaluation of Integral Transforms, Ph. D. Thesis, University of Utah, Salt Lake City, UT, 1978.
14. Lundin, L. and F. Stenger, Cardinal-Type Approximations of a Function and its Derivatives, SIAM J. Math. Anal., 10 (1979), pp. 139-160.
15. McNamee, J., F. Stenger and E. L. Whitney, Whittaker's Cardinal Function in Retrospect, Math. Comp., 25 (1971), pp. 141-154.
16. Mercier, B., Lectures on Topics of Finite Element Solutions of Elliptic Problems, Springer-Verlag, New York, 1979.
17. Meyers, N. and J. Serrin, $H=W$, Proc. Nat. Acad. Sci. USA, 51 (1964), pp. 1055-1056.
18. Mikhlin, S. G., Variational Methods in Mathematical Physics, Macmillan Co., New York, 1964.
19. Oden, J. T. and J. N. Reddy, An Introduction to the Mathematical Theory of Finite Elements, John Wiley and Sons, Inc., New York, 1976.
20. Schoombie, S. W. and J. F. Botha, Error Estimates for the Solution of the Radial Schrodinger Equation by the Rayleigh-Ritz Finite Element Method, IMA J. Numer. Anal., 1 (1981), pp. 47-63.
21. Schwing, J., Integral Equation Method of Solution of Potential Theory Problems in R^3 , Ph.D. Thesis, University of Utah, Salt Lake City, UT, 1976.
22. Shore, B. W., Solving the Radial Schrodinger Equation by Using Cubic-Spline Basis Functions, J. Chem. Phys., 58 (1973), pp. 3855-3866.
23. Shore, B. W., Comparison of Matrix Methods Applied to the Radial Schrodinger Eigenvalue Equation: The Morse Potential, J. Chem. Phys., 59 (1973), pp. 6450-6463.

24. Shore, B. W., Use of the Rayleigh-Ritz-Galerkin Method with Cubic-Splines for Constructing Single-Particle Bound-State Radial Wave Functions: The Hydrogen Atom and its Spectrum, J. Phys. B: Atomic and Molec. Phys., 6 (1973), pp. 1923-1932.
25. Stenger, F., Approximations via Whittaker's Cardinal Function, J. Approx. Thy., 17 (1976), pp. 222-240.
26. Stenger, F., A Sinc-Galerkin Method of Solution of Boundary Value Problems, Math. Comp., 33 (1979), pp. 85-109.
27. Stenger, F., Numerical Methods Based on Whittaker Cardinal, or Sinc Functions, SIAM Review, 23 (1981), pp. 165-224.
28. Steward, G. W., Introduction to Matrix Computations, Academic Press, New York, 1973.
29. Strang, G. and G. J. Fix, An Analysis of the Finite Element Method, Prentice-Hall, Inc., Englewood Cliffs, NJ, 1973.
30. Tadros, G. S. and A. Ghali, Convergence of Semianalytical Solution of Plates, J. Engrg. Mech. Div., ASCE 99 EM5 (1973), pp. 1023-1035.
31. Taylor, A. E. and D. C. Lay, Introduction to Functional Analysis, second edition, John Wiley and Sons, Inc., New York, 1980.
32. Wendroff, B., Bounds for Eigenvalues of some Differential Operators by the Rayleigh-Ritz Method, Math. Comp., 19 (1965), pp. 218-224.
33. Whittaker, E. T., On the Functions which are Represented by the Expansions of the Interpolation Theory, Proc. Roy. Soc. Edinburgh, 35 (1915), pp. 181-194.
34. Zlamal, M., The Finite Element Method in Domains with Curved Boundaries, Internat. J. Numer. Anal. Engrg., 5 (1973), pp. 367-373.
35. Zlamal, M., Curved Elements in the Finite Element Method I, SIAM J. Numer. Anal., 10 (1973), pp. 229-240.
36. Zlamal, M., Curved Elements in the Finite Element Method II, SIAM J. Numer. Anal., 11, (1974), pp. 347-362.

37. Zygmund, A., Trigonometric Series, Cambridge University Press, London, 1959.

MONTANA STATE UNIVERSITY LIBRARIES
stks D378.R452@Theses RL
Galerkin schemes for elliptic boundary v



3 1762 00113707 2

D378
R452
Cap. 2

