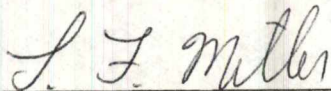


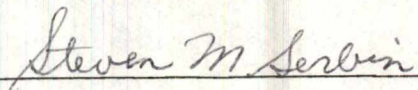
To the Graduate Council:

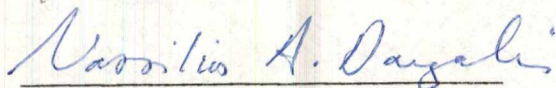
I am submitting herewith a dissertation written by Edward J. Allen entitled "A Galerkin Method for Numerically Solving the Energy-Dependent Neutron Transport Equation." I have examined the final copy of this dissertation for form and content and recommend that it be accepted in partial fulfillment of the requirements for the degree of Doctor of Philosophy, with a major in Mathematics.

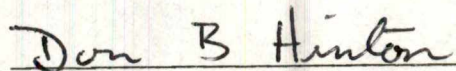

F.W. Stallmann, Major Professor

We have read this dissertation
and recommend its acceptance:

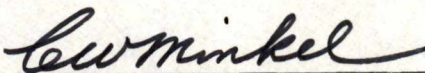








Accepted for the Council:


The Graduate School

A GALERKIN METHOD FOR NUMERICALLY SOLVING THE
ENERGY-DEPENDENT NEUTRON TRANSPORT EQUATION

A Dissertation
Presented for the
Doctor of Philosophy
Degree

The University of Tennessee, Knoxville

Edward James Allen

August 1983

DEDICATION

To Linda.

ACKNOWLEDGEMENTS

I am very grateful to Professor F.W. Stallmann for his guidance, encouragement, and assistance in this work. I very much appreciate the effort spent by the dissertation committee members for reading this dissertation. I would like to thank Sue Smith of the Computer Science Department for her assistance with the University's computer system. I am also very grateful to Cindi Blair for typing this dissertation. In addition, I am very grateful to my wife, Linda, whose encouragement and assistance made this work possible.

ABSTRACT

In the present investigation, Galerkin's method is applied to the first-order form of the steady-state energy-dependent neutron transport equation. The energy variable is treated separately from the direction and position variables yielding a set of multigroup equations which are a generalization of the standard multigroup equations. Numerical examples indicate greater accuracy is obtained using these generalized multigroup equations than using the standard multigroup equations. In addition, an error analysis is performed in which the energy variable is considered separately from the direction and position variables.

This method is applied to the one-dimensional plane, spherical, and cylindrical multigroup transport equations. Expansion functions and numerical integration rules are selected so that the inner products resulting in Galerkin's method are evaluated exactly and so that the flux density approximation is a continuous piecewise polynomial with respect to the position variable. In plane geometry, Galerkin's method using a piecewise linear approximation in the position variable leads to equations that are identical to the finite-difference discrete ordinates equations. In addition, using a piecewise parabolic approximation in Galerkin's method appears to give the same equations as a new fourth-order finite-difference method. It is demonstrated that these equations can be efficiently solved noniteratively rather than iteratively as is done presently. In spherical and cylindrical

geometries, Galerkin's method does not yield equations that are equivalent to the finite-difference discrete ordinates equations for any order of approximation. However, it is likely that the equations obtained using Galerkin's method for these geometries can be solved noniteratively in an efficient manner similar to that for plane geometry.

TABLE OF CONTENTS

CHAPTER	PAGE
I. INTRODUCTION.	1
II. THE FIRST-ORDER NEUTRON TRANSPORT EQUATION.	8
III. A GALERKIN METHOD	11
1. Introduction.	11
2. Theory.	15
3. Error Analysis.	24
IV. MULTIGROUP EQUATIONS.	36
1. Derivation.	37
2. Convergence Result.	44
V. PLANE GEOMETRY.	50
1. Derivation of Galerkin Equations.	52
2. Noniterative Method of Solution	63
VI. SPHERICAL AND CYLINDRICAL GEOMETRY.	72
1. Spherical Geometry.	73
2. Cylindrical Geometry.	81
VII. NUMERICAL RESULTS	92
1. Comparisons of Multigroup Approximations.	93
2. Comparisons of Approximations for Plane Geometry. . .	102
VIII. SUMMARY AND CONCLUSIONS	109
LIST OF REFERENCES.	112
VITA.	117

CHAPTER I

INTRODUCTION

The following questions are posed. Can the standard multigroup transport approximation to the energy-dependent neutron transport equation be improved by using a finite element approach? Can a finite element method for numerically solving the multigroup transport equations be formulated, specifically for one-dimensional geometries, such that flux densities at specified positions and directions can be efficiently calculated with greater accuracy than that obtained using well-established methods?

It will be shown in this dissertation that each of the above questions has an affirmative answer. Before proceeding, the status of current research relevant to these questions is summarized.

Numerical methods to solve the energy-dependent neutron transport equation has been an area of active research for some time. Methods that have undergone considerable research and development include: the finite-difference method, the finite element method, and the Monte Carlo method. These methods are generally applied to the multigroup transport equations which are derived by integrating the energy-dependent transport equation over specified energy intervals. However, the theory of multigroup transport with respect to numerical methods such as the finite element method is still under development by many researchers.¹ In addition, the development of numerical methods to accurately and efficiently approximate multigroup neutron flux densities with respect to position and direction variables is an area of active research.²

The neutron transport equation is solved numerically using deterministic methods or by using stochastic, i.e. Monte Carlo, methods. Monte Carlo methods are often used in transport calculations especially for three-dimensional shielding problems. However, Monte Carlo methods are not competitive with deterministic methods for one-dimensional problems.³ Deterministic methods can be classified, as in reference (2), according to whether the method uses the integrodifferential, integral, or surface integral form of the transport equation. All of these deterministic methods are applied to the one-group or multigroup transport equations.

The analysis presented in this paper is based on the integrodifferential form of the energy-dependent transport equation. Stochastic methods as well as methods based on other forms of the transport equation are not considered here.

One of the most widely used deterministic methods based on the integrodifferential form of the transport equation is the finite-difference discrete ordinates methods, especially for one- and two-dimensional geometries.^{4,5,6} The finite-difference discrete ordinates method has become a standard tool for neutronic work. It is also the method with which newer methods are generally compared. There is still much active research going on in developing methods which are more accurate or faster than the finite-difference discrete ordinates method, even for one-dimensional plane geometry.^{7,8,9,10,11,12,13} In the finite-difference discrete ordinates method, discrete directions, that is discrete values of direction cosines, are chosen. These discrete directions (discrete ordinates) correspond with points in a numerical

quadrature rule by which the integral term in the integrodifferential transport equation is approximated. A set of coupled differential equations are then obtained for the flux density as a function of position for each discrete direction. If a finite-difference approach is used for solving these equations, the finite-difference discrete ordinates equations are obtained.

Two remarks are now made concerning the finite-difference discrete ordinates method for later comparisons. The first remark is that the discrete ordinates equations are solved iteratively over each energy group rather than noniteratively. The second remark is that in spherical (and cylindrical) geometries supplementary equations are required. These equations relate the flux densities between adjacent direction and position mesh points. The well-known diamond-difference scheme assumes these relations to be linear.

Another deterministic method for solving the integrodifferential transport equation is the finite element method.^{1,2,14,15,16,17,18,19,20,21,22} Finite element methods can be divided into two groups depending on whether the method is based on the first-order form or second-order even-parity form of the integrodifferential transport equation. A summary of the advantages and disadvantages of basing a finite element method on the first-order equation or second-order equation is given in reference (16).

Numerical comparisons between discrete ordinates and finite element methods have been made.^{15,16,19} Finite element methods offer certain advantages over finite-difference methods especially for multidimensional geometries but it is stated in reference (16)

that the finite element method is not yet competitive with the highly optimized discrete ordinates codes.

All the deterministic methods described, whether discrete ordinates or finite element, are applied to one-group or multigroup transport calculations. Present deterministic methods therefore numerically solve the multigroup transport equations rather than the energy-dependent transport equation. The error obtained in approximating the energy-dependent neutron flux density by multigroup flux densities is generally overlooked in error analyses. In addition, as partial motivation for the present investigation, it is reasonable to suppose that the multigroup representation of the energy-dependent transport equation can be improved and deterministic methods applied to this improved form of multigroup equation.

There are two aims of the present investigation. The first aim is to derive, using Galerkin's method, a set of multigroup equations which are a generalization of the standard multigroup equations. It is desirable that these generalized multigroup equations provide a neutron flux density approximation with respect to the energy variable which is sufficiently more accurate than that obtained using the standard multigroup equations. The second aim is to formulate equations that relate the flux densities at specified positions and directions by applying Galerkin's method to the multigroup transport equations. It is desirable that these equations can be efficiently solved noniteratively over each energy group to provide flux density values of greater accuracy than those obtained using well-established methods.

In the present investigation, Galerkin's method is applied to the integrodifferential first-order form of the steady-state energy-dependent neutron transport equation. In this method, the flux density is expanded in linearly independent functions of energy, position, and direction. The dependency on energy is treated in some depth. The energy variable is treated separately from the position and direction variables and the expansion functions in energy are carefully chosen. This approach leads to a set of multigroup equations which are a generalization of the standard multigroup equations.

In addition to the analysis and derivation of a more general set of multigroup equations, Galerkin's method is applied to plane, spherical, and cylindrical one-dimensional multigroup transport equations. By suitably choosing the expansion functions in direction and position variables, equations that can be solved for flux densities at specified positions and directions are obtained. The expansion functions in the position variable are continuous piecewise polynomials. The expansion functions in the direction variables are selected in conjunction with numerical integration rules such that the inner products obtained in Galerkin's method are exactly evaluated. Although Galerkin's method is applied to the first-order multigroup transport equation in references (15), (16), (20), and (21), the analysis in the present investigation differs considerably in the expansion functions and numerical integration rules selected as well as in the derivation of equations relating flux densities at specified positions and directions.

Several interesting new results obtained in this work are described.

First, multigroup equations that are a generalization of the standard multigroup equations are derived using Galerkin's method. The standard multigroup equations are obtained if a single expansion function in energy over each energy interval is used in Galerkin's method. The accuracy of the generalized multigroup equations increases as the number of expansion functions in energy over each energy interval increases. Numerical results indicate that using two expansion functions in energy over each energy interval yields results as accurate as standard multigroup calculations with eight times the number of energy intervals.

Second, an error analysis is performed in which the energy variable is treated separately from the position and direction variables. The error in approximating the energy-dependent transport equation by multigroup equations is thus estimated.

Third, Galerkin's method applied to neutron transport in plane geometry results in, for a linear approximation in the position variable, the finite-difference discrete ordinates equations. This results in greater understanding of these particular discrete ordinates equations. Also, it is shown that these equations can be efficiently solved noniteratively over each energy group rather than iteratively as is done at present. It is also interesting that applying Galerkin's method with a parabolic approximation in the position variable appears to give the same equations as the new fourth-order finite-difference scheme for slab geometry described in reference (13). Higher order approximations in the position variable lead to equations of increasingly greater accuracy which can also be efficiently solved noniteratively.

Fourth, in spherical and cylindrical geometries, Galerkin's method does not lead to equations that are equivalent to the finite-difference discrete ordinates equations for any order of approximation. However, the equations derived using Galerkin's method are likely to provide more accurate results than the diamond-difference discrete ordinates equations as the number of discrete directions increases. The reason for this possible increase in accuracy is that in the diamond-difference discrete ordinates equations in spherical and cylindrical geometries it is assumed that the flux density is linear between adjacent discrete directions, while no such assumption is present in Galerkin's method.

The notation used in this paper is consistent when practicable with the notation used in the book by Bell and Glasstone.²³

CHAPTER II

THE FIRST-ORDER NEUTRON TRANSPORT EQUATION

The energy-dependent neutron transport equation is derived from neutron balance relations and describes the transport of neutrons in a scattering and absorbing medium with the presence of neutron sources. The dependent variable in the equation is the neutron flux density (also referred to as neutron flux) which is a function of neutron position, direction, and energy.

A one-group transport equation is derived from the energy-dependent transport equation by considering only neutrons within a given energy interval using appropriately defined group macroscopic cross sections and group neutron flux density. (The term group refers to an energy interval and macroscopic cross section refers to the probability per unit distance of a neutron interaction.) The multi-group equations are a coupled set of one-group equations. The "standard" multigroup equations, as defined in this paper, are the multigroup equations derived assuming separability of energy from position and direction variables (i.e. the neutron flux density in any energy interval is assumed to be the product of a function dependent on energy and a function dependent on position and direction).

The first-order steady-state energy-dependent neutron transport equation has the form²³:

$$\vec{\Omega} \cdot \nabla \psi(\vec{r}, \vec{\Omega}, E) + \sigma(\vec{r}, E) \psi(\vec{r}, \vec{\Omega}, E) = \quad (1)$$

$$\int_B \int_F \psi(\vec{r}, \vec{\Omega}', E') \sigma(\vec{r}, E') f(\vec{r}, \vec{\Omega}', E' \rightarrow \vec{\Omega}, E) d\vec{\Omega}' dE'$$

$$+ Q(\vec{r}, \vec{\Omega}, E)$$

where

$\vec{r}$ = position vector, $\vec{r} \in D$, where D is a bounded convex domain in $\mathbb{R}^3$ with a Lipschitz continuous boundary

$\vec{\Omega}$ = direction vector of unit length, $\vec{\Omega} \in F$ where

$$F = \{\vec{\Omega}: \vec{\Omega} = (\Omega_1, \Omega_2, \Omega_3) \in \mathbb{R}^3 \text{ and } |\vec{\Omega}| = 1\}$$

E = energy, $E \in B$ where $B = \{E \in \mathbb{R}^1\}$

$\psi(\vec{r}, \vec{\Omega}, E)$ = neutron flux density of energy E and direction $\vec{\Omega}$ at position $\vec{r}$

$\sigma(\vec{r}, E)$ = total macroscopic cross section at position $\vec{r}$ for neutrons of energy E ; probability per unit distance of a neutron interaction at position $\vec{r}$ for neutrons of energy E

$\sigma(\vec{r}, E') f(\vec{r}; \vec{\Omega}', E' \rightarrow \vec{\Omega}, E) = \sigma f(\vec{r}; \vec{\Omega}', E' \rightarrow \vec{\Omega}, E) =$
probability per unit distance of a neutron transfer from $(\vec{\Omega}', E')$ to $(\vec{\Omega}, E)$ at position $\vec{r}$

$Q(\vec{r}, \vec{\Omega}, E)$ = neutron source at position $\vec{r}$ of energy E and direction $\vec{\Omega}$

$A = DXFXB$ where $(\vec{r}, \vec{\Omega}, E) \in A$

$d\vec{\Omega}$ = surface element on the unit sphere $\int_F d\vec{\Omega} = 4\pi$.

Generally, the energy spectrum (the functional dependence on energy) of the flux density is roughly known for all $(\vec{r}, \vec{\Omega}) \in DXF$. Assume that the energy spectrum of the flux density is nearly independent of $\vec{r}$ and $\vec{\Omega}$ for all $(\vec{r}, \vec{\Omega}) \in DXF$ and is approximately $b(E)$ where $b(E)$

is piecewise continuous in B and $0 < b_{\min} \leq b(E) \leq b_{\max} < \infty$ for all $E \in B$. For example, $b(E)$ may be the familiar Maxwellian - $1/E$ - fission spectrum. Now let $\psi(\vec{r}, \vec{\Omega}, E) = b(E) \phi(\vec{r}, \vec{\Omega}, E)$. It is to be expected that $\phi(\vec{r}, \vec{\Omega}, E)$ is only slowly varying with respect to energy.

With this substitution equation (1) becomes:

$$\begin{aligned} \vec{\Omega} \cdot \nabla \phi(\vec{r}, \vec{\Omega}, E) b(E) + \sigma(\vec{r}, E) b(E) \phi(\vec{r}, \vec{\Omega}, E) = & \quad (2) \\ = \int_B \int_F b(E') \phi(\vec{r}, \vec{\Omega}', E') \sigma f(\vec{r}; \vec{\Omega}', E' \rightarrow \vec{\Omega}, E) d\vec{\Omega}' dE' \\ & + Q(\vec{r}, \vec{\Omega}, E) \end{aligned}$$

As will be seen, the presence of $b(E)$ in equation (2) influences many of the results in Chapters III and IV.

The existence of a solution to equation (2) for a critical system or for a subcritical system with a neutron source is assumed in this paper based on physical grounds. Existence of a solution to the time-independent transport equation seems to have been rigorously proved only for special cases.²³

CHAPTER III

A GALERKIN METHOD

1. Introduction

Much of this section consists of defining notation to be used in later sections. Also, the weak form of equation (2) is introduced. It is the weak form of equation (2) for which an approximate solution is eventually sought.

In this report, $H^{1,1,0}(A)$ is a generalized Sobolev space.²⁴ The numbers 1,1,0 mean that functions in $H^{1,1,0}(A)$ are elements of $L^2(A)$ and have first order weak derivatives with respect to the position coordinates and direction coordinates. Explicitly, $H^{1,1,0}(A) = \{f(\vec{r}, \vec{\Omega}, E): f \in L^2(A), \frac{\partial f}{\partial x_i} \in L^2(A) \text{ for } i = 1, 2, \dots, N_x \text{ where } \frac{\partial f}{\partial x_i} \text{ is the weak derivative of } f \text{ with respect to } x_i \text{ and } N_x \text{ are the number of coordinates of } \vec{r} \text{ and } \vec{\Omega} \text{ or } (\vec{r}, \vec{\Omega}) = (x_1, x_2, \dots, x_{N_x})\}$.

The generalized Sobolev space, $H^{1,1,0}(A)$, is used in this report so that approximations to the neutron flux density may be discontinuous with respect to the energy variable. Errors in the approximation with respect to the energy variable, estimated in the third chapter of this chapter, then approach zero as the number of energy intervals approaches infinity.

The standard Sobolev spaces are denoted by $H^k(A)$ and

$$H^k(A) = \{f(\vec{r}, \vec{\Omega}, E): f \in L^2(A), D^\alpha f \in L^2(A) \text{ for } |\alpha| \leq k$$

where $D^\alpha f$ are the weak derivatives of $f\}$.

The corresponding norms have the forms:

$$\|f\|_{1,1,0}^2 = \|f\|^2 + \sum_{i=1}^{N_x} \left\| \frac{\partial f}{\partial x_i} \right\|^2$$

$$\|f\|_k^2 = \sum_{|\alpha| \leq k} \|D^\alpha f\|^2$$

where $\|f\|$ is the $L^2(A)$ -norm of f or

$$\|f\|^2 = \int_B \int_F \int_D |f|^2 d\vec{r} d\vec{\Omega} dE .$$

Note that

$$\|f\|_1^2 = \|f\|_{1,1,0}^2 + \left\| \frac{\partial f}{\partial E} \right\|^2$$

and

$$\|f\| \leq \|f\|_{1,1,0} .$$

For the purpose of the Galerkin method an inner product is defined as

$$(u, v) = \int_A u(\vec{r}, \vec{\Omega}, E) v(\vec{r}, \vec{\Omega}, E) d\vec{r} d\vec{\Omega} dE ;$$

thus

$$(u, u) = \|u\|^2$$

and by the Schwarz inequality, $|(u, v)| \leq \|u\| \|v\|$.

The weak formulation of the problem then takes the following form:

$$(\vec{\Omega} \cdot \nabla \phi, u) + (K\phi, u) = (Q, u) \quad \text{for every } u \in H^{1,1,0}(A) \quad (3)$$

$$\text{with } K\phi = \sigma(\vec{r}, E) \phi(\vec{r}, \vec{\Omega}, E) -$$

$$\int_B \int_F b(E') \phi(\vec{r}, \vec{\Omega}', E') \sigma(\vec{r}; \vec{\Omega}', E' \rightarrow \vec{\Omega}, E) d\vec{\Omega}' dE' .$$

Considered in the present investigation are two boundary conditions, which are most frequently encountered in neutron transport problems, namely prescribed inward flux densities and reflective boundary conditions. To describe these boundary conditions, appropriate notation has to be defined. Let

$$\delta D = \text{surface of } D$$

and

$$\Gamma = \{(\vec{r}, \vec{\Omega}, E) \in A: \vec{r} \in \delta D\} .$$

The boundary Γ is subdivided into two subsets Γ_1 and Γ_2 according to the two boundary conditions such that $\Gamma = \Gamma_1 \cup \Gamma_2$ and Γ_1 and Γ_2 are disjoint. The boundaries Γ_1 and Γ_2 are further subdivided to represent incoming and outgoing flux densities, i.e.,

$$\Gamma_{1+} = \{(\vec{r}, \vec{\Omega}, E) \in \Gamma_1: \vec{\Omega} \cdot n(\vec{r}) > 0\} \quad (4)$$

$$\Gamma_{1-} = \{(\vec{r}, \vec{\Omega}, E) \in \Gamma_1: \vec{\Omega} \cdot n(\vec{r}) \leq 0\}$$

and similarly for Γ_{2-} and Γ_{2+} where $n(\vec{r})$ is the outward normal at $\vec{r}$ on δD . Thus the two boundary conditions are defined as follows:

$$\begin{aligned} \phi(\vec{r}, \vec{\Omega}, E) &= \rho(\vec{r}, \vec{\Omega}, E) \text{ on } \Gamma_{1-} \text{ where} \\ b(E) \rho(\vec{r}, \vec{\Omega}, E) &\text{ is a given inward flux density,} \end{aligned} \quad (5)$$

$$\begin{aligned} \phi(\vec{r}, \vec{\Omega}, E) &= \phi(\vec{r}, M(\vec{\Omega}, n(\vec{r})), E) \text{ on } \Gamma_{2-} \text{ where} \\ M(\vec{\Omega}, n(\vec{r})) &\text{ is the mirror image of } \vec{\Omega} \text{ with respect} \\ &\text{to the surface normal } n(\vec{r}). \end{aligned} \quad (6)$$

Thus, on Γ_1 an inward flux density is specified and on Γ_2 a reflecting boundary is specified.

It is possible to select a function $y \in H^{1,1,0}(A)$ that satisfies the boundary conditions on Γ_1 and Γ_2 . Then the function defined by $\phi - y$ equals zero on Γ_{1-} and satisfies an equation identical in form to equation (3) but with a modified source term. Therefore, without loss of generality, the function ρ will be defined to be zero in the remainder of this paper.

Therefore, the weak formulation of the problem is to find $\phi \in H^{1,1,0}(A)$ such that

$$(\vec{\Omega} \cdot \nabla \phi b, u) + (K\phi, u) = (q, u) \text{ for every } u \in H^{1,1,0}(A) \quad (7)$$

$$\begin{cases} \phi(\vec{r}, \vec{\Omega}, E) = 0 & \text{on } \Gamma_{1-} \\ \phi(\vec{r}, \vec{\Omega}, E) = \phi(\vec{r}, M(\vec{\Omega}, n(\vec{r})), E) & \text{on } \Gamma_2 \end{cases} \quad (8)$$

where $q(\vec{r}, \vec{\Omega}, E) \in L^2(A)$ is a specified neutron source.

Equations (7) and (8) are the equations studied in the remainder of Chapter III. A function $\phi \in H^{1,1,0}(A)$ that satisfies equations (7) and (8) is the weak solution to equation (2). Actually, the neutron flux density, ψ , is of more interest than ϕ in this investigation and it can readily be shown that $\psi \in H^{1,1,0}(A)$ if and only if $\phi \in H^{1,1,0}(A)$ since $b \in \text{piecewise } C(B)$.

In the present investigation, an approximation to $\phi(\vec{r}, \vec{\Omega}, E) \in H^{1,1,0}(A)$ is sought. The approximation to $\psi(\vec{r}, \vec{\Omega}, E)$ is then the product of $b(E)$ and the approximation to $\phi(\vec{r}, \vec{\Omega}, E)$.

2. Theory

Many of the results in this section are based on the work of Ukai.²² However the analysis presented here is somewhat more general since an energy dependent function, $b(E)$, is included in the analysis and reflecting boundary conditions are allowed. Also the notation used here is consistent with that of Bell and Glasstone²³ which appears to be more widely adopted.

The main result of this section concerns the convergence of ϕ^h to ϕ where ϕ^h is an approximation to ϕ . The proof of this result relies on the assumption that there exist constants $s \in [0, 1)$ and $\gamma > 0$ such that $s(\vec{\Omega} \cdot \nabla v, v) + (Kv, v) \geq \gamma \|v\|^2$ for every $v \in H^{1,1,0}(A)$ that satisfies the boundary conditions given by equation (8). Sufficient conditions for this assumption to be valid are given. The first proposition of this section states that this assumption implies subcriticality of the system (although this assumption is not required for subcriticality). Before the convergence result is presented, some uniqueness and existence results are given.

$$\text{Define } Iv = \int_B \int_F b(E') v(\vec{r}, \vec{\Omega}', E') \sigma(\vec{r}; \vec{\Omega}', E' \rightarrow \vec{\Omega}, E) d\vec{\Omega}' dE'.$$

Based on physical considerations, it is assumed that I is a bounded operator on $L^2(A)$ and thus there exists a $c > 0$ such that

$$\|Iv\| \leq c\|v\| \text{ for every } v \in L^2(A). \text{ Assume that}$$

$0 \leq \sigma_{\min} \leq \sigma(\vec{r}, E) \leq \sigma_{\max} < \infty$. Since $Kv = \sigma(\vec{r}, E) b(E) v(\vec{r}, \vec{\Omega}, E) - Iv$, then

$$\|K\| = \sup_{v \in L^2(A)} \frac{\|Kv\|}{\|v\|} \leq \sup_{v \in L^2(A)} \frac{(\sigma_{\max} b_{\max} + c)\|v\|}{\|v\|} = \sigma_{\max} b_{\max} + c.$$

Thus K is a bounded operator on $L^2(A)$.

Consider $(\vec{\Omega} \cdot \nabla v, w)$. By Green's formula,^{25,26}

$$(\vec{\Omega} \cdot \nabla v, w) = -(bv, \vec{\Omega} \cdot \nabla w) +$$

$$\int_{\Gamma} b(E) v(\vec{r}, \vec{\Omega}, E) w(\vec{r}, \vec{\Omega}, E) (n(\vec{r}) \cdot \vec{\Omega}) dS d\vec{\Omega} dE$$

for every $v, w \in H^{1,1,0}(A)$ where dS is the surface element of δD .

If v also satisfies the boundary conditions given by equation (8), then

$$(\vec{\Omega} \cdot \nabla v, v) = \tag{9}$$

$$= \frac{1}{2} \int_{\Gamma_{1+}} b(E) v^2(\vec{r}, \vec{\Omega}, E) (n(\vec{r}) \cdot \vec{\Omega}) dS d\vec{\Omega} dE \geq 0,$$

since the integral over Γ_{1-} is zero and the integrals over Γ_{2-} and Γ_{2+} cancel each other.

Finally consider $(\vec{\Omega} \cdot \nabla v b, v) + (Kv, v)$. If v satisfies the boundary conditions on Γ_1 and on Γ_2 then by equation (9),

$$\begin{aligned} (\vec{\Omega} \cdot \nabla v b, v) + (Kv, v) &\geq (Kv, v) \\ &= (\sigma b v, v) - (Iv, v) \\ &\geq (\sigma_{\min} b_{\min} - c) \|v\|^2. \end{aligned}$$

In the remainder of this chapter it is assumed that there exists a $\gamma > 0$ and a $s \in [0, 1)$ such that

$$s(\vec{\Omega} \cdot \nabla v b, v) + (Kv, v) \geq \gamma \|v\|^2 \quad (10)$$

for every $v \in H^{1,1,0}(A)$ that satisfies the boundary conditions on Γ_1 and Γ_2 given by equation (8). A sufficient condition for equation (10) to hold is that $\sigma_{\min} b_{\min} > c$ in which case $s = 0$.

Notice that equations (9) and (10) imply that $(\vec{\Omega} \cdot \nabla v b, v) + (Kv, v) \geq \gamma \|v\|^2$ for every $v \in H^{1,1,0}(A)$ that satisfies equation (8). A value of s less than one is necessary in the proof of Theorem 1.

The criticality problem is considered to show that the assumption given by equation (10) implies subcriticality. The time-dependent transport equation has the form:²³

$$\frac{b}{v} \frac{\partial \phi}{\partial t} = -\vec{\Omega} \cdot \nabla \phi b - K\phi$$

where $\phi = \phi(\vec{r}, \vec{\Omega}, E, t)$

v = neutron speed

and t = time .

Consider solutions to the above equation of the form

$\phi(\vec{r}, \vec{\Omega}, E, t) = \phi(\vec{r}, \vec{\Omega}, E) e^{\alpha t}$. This substitution results in the following eigenvalue equation:

$$\frac{b}{v} \alpha_j \phi_j = -\vec{\Omega} \cdot \nabla \phi_j b - K \phi_j .$$

The weak form of the above equation is

$$\alpha_j \left(\frac{b}{v} \phi_j, u \right) = -(\vec{\Omega} \cdot \nabla \phi_j b, u) - (K \phi_j, u) \text{ for every } u \in H^{1,1,0}(A).$$

If α_0 is the value of α_j having the largest real part, then $\phi(\vec{r}, \vec{\Omega}, E, t)$ eventually varies as $e^{\alpha_0 t}$. Physically, α_0 is real since ϕ is everywhere real and nonnegative. If α_0 is less than zero, equal to zero, or greater than zero, the system is called subcritical, critical, or supercritical, respectively. If a system is critical or supercritical, a nonzero solution may exist to the above equation for $\alpha_j = 0$ in which case a unique solution does not exist to equation (7). However, as is shown in the following proposition, if equation (10) holds then the system is subcritical. Unfortunately, since many practical problems involve subcritical systems, subcriticality does not imply equation (10).

Proposition 1. Suppose $\phi \in H^{1,1,0}(A)$ is a solution of the eigenvalue equation with eigenvalue α_0 and equation (10) holds. Then $\alpha_0 < 0$ and the system is subcritical.

Proof: Assume otherwise, i.e. that the system is critical or supercritical and the eigenvalue α_0 is greater than or equal to zero.

Let $u = \phi$ in the eigenvalue equation where ϕ is not identically equal to zero. Hence, $(\vec{\Omega} \cdot \nabla \phi b, \phi) + (K\phi, \phi) = -\alpha_0 \left(\frac{b}{v}\phi, \phi\right)$. By equation (10), $0 < \gamma \|\phi\|^2 \leq -\alpha_0 \left(\frac{b}{v}\phi, \phi\right)$.

Since $\alpha_0, b, v \geq 0$, the right hand side of the above equation is less than or equal to zero. This is a contradiction which implies that the system is neither critical nor supercritical. Thus, equation (10) implies subcriticality. /

Note that in the remainder of this paper ϕ is time-independent (steady state solution).

Uniqueness for solutions of equations (7) and (8) follows from Proposition 2 presented below.

Proposition 2. Suppose a solution $\phi \in H^{1,1,0}(A)$ of equations (7) and (8) exists and that there exists a $\gamma > 0$ such that equation (10) holds, then $\|\phi\| \leq \frac{1}{\gamma} \|q\|$.

Proof: Let $u = \phi$ in equation (7). Then $(\vec{\Omega} \cdot \nabla \phi b, \phi) + (K\phi, \phi) = (q, \phi)$. By equation (10), $\gamma \|\phi\|^2 \leq (q, \phi) \leq \|q\| \|\phi\|$. Hence, $\gamma \|\phi\| \leq \|q\|$. /

It is easy to see that Proposition 2 implies uniqueness of solutions of equations (7) and (8). Let ϕ_1 and ϕ_2 both be solutions of equations (7) and (8). Then $\phi_1 - \phi_2$ satisfies equations (7) and (8) with a source equal to zero. Hence by Proposition 2, $\|\phi_1 - \phi_2\| \leq 0$ implying that $\phi_1 = \phi_2$. Thus, if equation (10) holds, the solution of equations (7) and (8) is unique.

Approximations to ϕ in a finite dimensional space S^h are studied, where $S^h \subset H^{1,1,0}(A)$. Let

$$S^h = \{v^h \in H^{1,1,0}(A): v^h(\vec{r}, \vec{\Omega}, E) = \quad (11)$$

$$= \sum_{m=1}^M \sum_{g=1}^G \sum_{j=1}^J \xi_{jmg} w_{jmg}(\vec{r}, \vec{\Omega}, E)$$

where each w_{jmg} satisfies the boundary conditions and

$$w_{jmg}(\vec{r}, \vec{\Omega}, E) = \gamma_j(\vec{r}, \vec{\Omega}) b_{mg}(E) \chi_g(E).$$

Thus every $w \in S^h$ satisfies the conditions (8).

In the above expression for the set S^h , $\gamma_j(\vec{r}, \vec{\Omega}) b_{mg}(E) \chi_g(E) \in H^{1,1,0}(A)$ are linearly independent functions and

$$\chi_g(E) = \begin{cases} 1 & \text{if } E_{g+1} \leq E \leq E_g \\ 0 & \text{otherwise} \end{cases}$$

where a partition, Δ , of $B = \{E \in \mathbb{R}^1: 0 \leq E_{\min} \leq E \leq E_{\max} < \infty\}$ is defined as follows:

$$\Delta = \{E_g \in B, g = 1, 2, \dots, G+1:$$

$$E_{\max} = E_1 \geq E_2 \geq \dots \geq E_{G+1} = E_{\min}\}.$$

By convention, $E_s \leq E_t$ if $s > t$.

Let $\phi^h \in S^h$ and suppose that ϕ^h satisfies equation (12).

Then, as will be seen, ϕ^h is an approximation to ϕ .

$$(\vec{\Omega} \cdot \nabla \phi^h, w) + (K\phi^h, w) = (q, w) \quad \text{for every } w \in S^h. \quad (12)$$

Before a convergence proof is presented, existence and uniqueness of ϕ^h is considered.

If $\phi^h \in S^h$ satisfies equation (12), then, by a proof analogous to that of Proposition 2, $\|\phi^h\| \leq \frac{1}{\gamma} \|q\|$. Since S^h is finite dimensional, this inequality establishes the existence and uniqueness of ϕ^h .

In Theorem 1, the main result of this section is stated and proved concerning the convergence of ϕ^h to ϕ . For the proof, the inequality $\|\vec{\Omega} \cdot \nabla v\| \leq \|v\|_{1,1,0}$ for $v \in H^{1,1,0}(A)$ is needed. This inequality follows from Minkowski's inequality and from the fact that $|\Omega_i| \leq 1$ for $i = 1, 2, 3$.

Theorem 1. Suppose $\phi \in H^{1,1,0}(A)$ is the solution of equations (7) and (8) and there exists a $\gamma > 0$ and a $s \in [0, 1)$ such that equation (10) holds. Also suppose that $q \in L^2(A)$ and $\phi^h \in S^h$ is the solution of equation (12). Then there exists a $c_0 > 0$ such that

$$\|\phi - \phi^h\| \leq c_0 \inf_{u^h \in S^h} \|\phi - u^h\|_{1,1,0}. \quad (13)$$

Proof: Notice that

$$\begin{aligned} (b\vec{\Omega} \cdot \nabla(\phi - \phi^h), \phi - \phi^h) + (K(\phi - \phi^h), \phi - \phi^h) = \\ (b\vec{\Omega} \cdot \nabla(\phi - \phi^h), \phi - u^h) + (K(\phi - \phi^h), \phi - u^h) \end{aligned}$$

since $(b\vec{\Omega} \cdot \nabla(\phi - \phi^h), w^h) + (K(\phi - \phi^h), w^h) = 0$ for every $w^h \in S^h$. Inequality (9) applied to $v = (1 - s)^{\frac{1}{2}}(\phi - \phi^h) - (1 - s)^{-\frac{1}{2}}(\phi - u^h)$ yields

$$\begin{aligned} (b\vec{\Omega} \cdot \nabla(\phi - \phi^h), \phi - u^h) &\leq (1 - s)(b\vec{\Omega} \cdot \nabla(\phi - \phi^h), \phi - \phi^h) \\ &+ \frac{1}{1 - s} (b\vec{\Omega} \cdot \nabla(\phi - u^h), \phi - u^h) - (b\vec{\Omega} \cdot \nabla(\phi - u^h), \phi - \phi^h) \end{aligned}$$

Applying this inequality to the above equality and subtracting the term $(1 - s)(b\vec{\Omega} \cdot \nabla(\phi - \phi^h), \phi - \phi^h)$ on both sides yields

$$\begin{aligned} s(b\vec{\Omega} \cdot \nabla(\phi - \phi^h), \phi - \phi^h) + (K(\phi - \phi^h), \phi - \phi^h) &\leq \\ (K(\phi - \phi^h), \phi - u^h) + \frac{1}{1 - s} (b\vec{\Omega} \cdot \nabla(\phi - u^h), \phi - u^h) \\ - (b\vec{\Omega} \cdot \nabla(\phi - u^h), \phi - \phi^h) . \end{aligned}$$

The left hand side of the above inequality is greater than or equal to $\gamma \|\phi - \phi^h\|^2$ according to inequality (10). Hence,

$$\begin{aligned} \gamma \|\phi - \phi^h\|^2 &\leq \|K\| \|\phi - \phi^h\| \|\phi - u^h\| + \frac{b_{\max}}{1 - s} \|\phi - u^h\|_{1,1,0} \|\phi - u^h\| \\ &+ b_{\max} \|\phi - u^h\|_{1,1,0} \|\phi - \phi^h\| \\ &\leq (b_{\max} + \|K\|) \|\phi - u^h\|_{1,1,0} \|\phi - \phi^h\| \\ &+ \frac{b_{\max}}{1 - s} \|\phi - u^h\|_{1,1,0}^2 . \end{aligned}$$

Setting $R = \frac{\|\phi - u^h\|_{1,1,0}}{\|\phi - \phi^h\|}$ we obtain

$$\gamma \leq (b_{\max} + \|K\|) R + \frac{b_{\max}}{1-s} R^2.$$

Noting that $R > 0$ it follows that $R \geq \frac{1}{c_0}$ where $\frac{1}{c_0}$ is the positive root of the quadratic equation $\frac{b_{\max}}{1-s} x^2 + (b_{\max} + \|K\|)x - \gamma = 0.$

In the finite element method, the sets S^h are defined in such a manner that $\inf_{u^h \in S^h} \|f - u^h\|_{1,1,0}$ goes to zero for $h \rightarrow 0$ for any $f \in H^{1,1,0}$. If the sets S^h are so constructed, then the above result implies convergence of ϕ^h to ϕ in the $L^2(A)$ -norm.

Also note that $\|\psi - \psi^h\| \leq b_{\max} \|\phi - \phi^h\|$ so as ϕ^h converges to ϕ , $\psi^h = b\phi^h$ converges to ψ .

Before continuing, one remark is made concerning the convergence result given in Theorem 1. In Theorem 1, the elements of S^h satisfy the boundary conditions. However, in Chapters V and VI it is convenient to require the approximation, ϕ^h , to satisfy the boundary conditions only at specified nodal points, i.e. only at discrete directions on δD . This type of approximation is common in finite element work and the reader is referred to references (27) and (28). The proof given in Theorem 1 can be modified to include this boundary approximation but the calculation is tedious and is not presented.

3. Error Analysis

In the last section, the estimate $\|\phi - \phi^h\| \leq c_0 \inf_{u^h \in S^h} \|\phi - u^h\|_{1,1,0}$

was obtained where $\phi^h \in S^h$ is an approximation to $\phi \in H^{1,1,0}(A)$. It is the aim of this section to show that if S^h is appropriately constructed ϕ^h does converge to ϕ as the number of basis elements in S^h increases. Analyses of this type are generally performed by choosing $v^h \in S^h$ that interpolates ϕ and showing that the interpolation error $\|\phi - v^h\|_{1,1,0}$ approaches zero as the number of basis elements in S^h approaches infinity.

This analysis has been performed in several publications,^{16,20,21} but only in respect to the position and direction variables, disregarding the approximation in the energy variable. The present investigation attempts to fill this gap by discussing the approximation in the energy variable analyzed separately from the position and direction variables. The approximation error in the position and direction variables is not given here explicitly since it is discussed sufficiently in the literature.

Due to the complexity of the analysis, only two cases are considered, $M = 1$ and $M = 2$. For $M = 1$, a step function interpolation in the energy variable is used. The results indicate that the approximation error is proportional to the maximum energy interval width. This is an interesting result as it will be shown in the next chapter that $M = 1$ corresponds to the standard multigroup equations. For $M = 2$, a linear interpolation is used in the energy variable. The results indicate that the approximation error is proportional to

the square of the maximum energy interval width. This result indicates that greater accuracy should be expected using $M = 2$ than using $M = 1$. In order to facilitate the error estimates, it is assumed that $\phi \in H^2(A)$ or $\phi \in H^3(A)$ respectively for $M = 1$ and $M = 2$. These assumptions are reasonable on physical grounds. It may be possible to relax these assumptions somewhat, but at the cost of greatly complicating the analysis.

A function z is selected, which satisfies the boundary conditions given by equation (8), and has the following form:

$$z(\vec{r}, \vec{\Omega}, E) = \sum_{m=1}^M \sum_{g=1}^G \chi_g(E) b_{mg}(E) \phi_{mg}(\vec{r}, \vec{\Omega})$$

$$\text{where } \phi_{mg}(\vec{r}, \vec{\Omega}) \in H^1(\text{DXF}) .$$

Also, recall that $u^h \in S^h$ has the form:

$$u^h(\vec{r}, \vec{\Omega}) = \sum_{m=1}^M \sum_{g=1}^G \chi_g(E) b_{mg}(E) u_{mg}^h(\vec{r}, \vec{\Omega})$$

$$\text{where } u_{mg}^h = \sum_{j=1}^J \gamma_j(\vec{r}, \vec{\Omega}) \xi_{jmg} .$$

Therefore, from Theorem 1,

$$\|\phi - \phi^h\| \leq c_0 \inf_{u^h \in S^h} \|\phi - u^h\|_{1,1,0}$$

$$\leq c_0 \|\phi - z\|_{1,1,0} + c_0 \inf_{u^h \in S^h} \|z - u^h\|_{1,1,0} . \quad (14)$$

In this section, the two terms on the right hand side of equation (14) are estimated. If z is properly chosen, the first term on the right hand side provides an estimate of the approximation error with respect to the energy variable. The second term can then provide an estimate of the approximation error with respect to the position and direction variables. The error analysis is thus separated into two steps. The first step is to estimate $\|\phi - z\|_{1,1,0}$ where $\phi \in H^{1,1,0}(A)$ and z is selected to make the error small. The second step is to estimate $\inf_{u^h \in S^h} \|z - u^h\|_{1,1,0}$.

A. Step 1 - Analysis of $\|\phi - z\|_{1,1,0}$.

It is difficult to estimate $\|\phi - z\|_{1,1,0}$ for general functions $b_{mg}(E)$. However, two cases are of particular interest. These cases are $M = 1$ and $M = 2$ and $b_{1g}(E)$ and $b_{2g}(E)$ are set equal to 1 and E , respectively.

a. Case 1: $M = 1$. Recall that $\psi(\vec{r}, \vec{\Omega}, E) = b(E) \phi(\vec{r}, \vec{\Omega}, E)$ where $b(E)$ approximates the energy distribution of the flux density $\psi(\vec{r}, \vec{\Omega}, E)$. Thus, $\phi(\vec{r}, \vec{\Omega}, E)$ may not vary greatly with respect to the energy variable E . Therefore, for $M = 1$, it is logical to let $b_{1g}(E) = 1.0$ for $1 \leq g \leq G$. In fact, this is reasonable even for $M \geq 1$ and throughout the remainder of this report it will be assumed that $b_{1g}(E) = 1.0$ for $1 \leq g \leq G$.

As mentioned above we assume that $\phi(\vec{r}, \vec{\Omega}, E) \in H^2(A) \subset H^{1,1,0}(A)$.

With the above assumption, for almost every $(\vec{r}, \vec{\Omega}) \in D \times F$, $\phi(\vec{r}, \vec{\Omega}, E)$ or any first order weak derivative of $\phi(\vec{r}, \vec{\Omega}, E)$ is equivalent to a function which is absolutely continuous on B . This can be seen in the following manner.

Let $\frac{\partial \phi}{\partial x}$ be a (weak) L^2 -derivative of ϕ with respect to x where x is a coordinate of $\vec{r}$, $\vec{\Omega}$, or E . Then since $\phi \in H^2(A)$, $\frac{\partial \phi}{\partial x} \in L^2(B)$ for almost every $(\vec{r}, \vec{\Omega}) \in D \times F$ and $\frac{\partial \phi}{\partial x}$ has a (weak) L^2 -derivative $\frac{\partial^2 \phi}{\partial E \partial x} \in L^2(B)$ for almost every $(\vec{r}, \vec{\Omega}) \in D \times F$. Therefore, for almost every $(\vec{r}, \vec{\Omega}) \in D \times F$, $\frac{\partial \phi}{\partial x}$ is equivalent to a function which is absolutely continuous on B by Proposition 1.2 of reference (29).

Theorem 2. Suppose $M = 1$, $\phi \in H^2(A)$, $b_{1g}(E) = 1.0$ for $1 \leq g \leq G$ and $z = \sum_{g=1}^G \chi_g(E) \phi(\vec{r}, \vec{\Omega}, E_{g+1})$, i.e. $\phi_{1g}(\vec{r}, \vec{\Omega}) = \phi(\vec{r}, \vec{\Omega}, E_{g+1})$. Then

$$\|\phi - z\|_{1,1,0} \leq \max_{1 \leq g \leq G} \frac{E_g - E_{g+1}}{\sqrt{2}} \left\| \frac{\partial \phi}{\partial E} \right\|_1.$$

Proof: Note that

$$\|\phi - z\|_{1,1,0}^2 = \left\| \sum_{g=1}^G \chi_g(E) (\phi(\vec{r}, \vec{\Omega}, E) - \phi(\vec{r}, \vec{\Omega}, E_{g+1})) \right\|_{1,1,0}^2 \quad (15)$$

$$\begin{aligned} &= \left\| \sum_{g=1}^G \chi_g(E) (\phi(\vec{r}, \vec{\Omega}, E) - \phi(\vec{r}, \vec{\Omega}, E_{g+1})) \right\|^2 \\ &+ \sum_{i=1}^N \left\| \sum_{g=1}^G \chi_g(E) \left(\frac{\partial \phi}{\partial x_i}(\vec{r}, \vec{\Omega}, E) - \frac{\partial \phi}{\partial x_i}(\vec{r}, \vec{\Omega}, E_{g+1}) \right) \right\|^2 \end{aligned}$$

where $(\vec{r}, \vec{\Omega}) = (x_1, x_2, \dots, x_{N_x})$

and $\frac{\partial}{\partial x_i}$ is the (weak) L^2 -derivative with respect to x_i .

Since $\phi \in H^2(A)$, ϕ or any first order weak derivative of ϕ is equivalent to a function which is absolutely continuous with respect to E . Hence,

$$\phi(\vec{r}, \vec{\Omega}, E) - \phi(\vec{r}, \vec{\Omega}, E_{g+1}) = \int_{E_{g+1}}^E \frac{\partial \phi}{\partial E'}(\vec{r}, \vec{\Omega}, E') dE' \quad (16)$$

and

$$\begin{aligned} \frac{\partial \phi}{\partial x_i}(\vec{r}, \vec{\Omega}, E) - \frac{\partial \phi}{\partial x_i}(\vec{r}, \vec{\Omega}, E_{g+1}) &= \\ &= \int_{E_{g+1}}^E \frac{\partial^2 \phi}{\partial E' \partial x_i}(\vec{r}, \vec{\Omega}, E') dE'. \end{aligned} \quad (17)$$

Thus,

$$\begin{aligned} \|\phi - z\|_{1,1,0}^2 &= \left\| \sum_{g=1}^G \chi_g(E) \int_{E_{g+1}}^E \frac{\partial \phi}{\partial E'}(\vec{r}, \vec{\Omega}, E') dE' \right\|^2 \\ &+ \sum_{i=1}^N \left\| \sum_{g=1}^G \chi_g(E) \int_{E_{g+1}}^E \frac{\partial^2 \phi}{\partial E' \partial x_i}(\vec{r}, \vec{\Omega}, E') dE' \right\|^2. \end{aligned} \quad (18)$$

Consider just one of the norms on the right hand side of equation (18). Any of them can be evaluated in a similar manner. Consider

$$\begin{aligned} &\left\| \sum_{g=1}^G \chi_g(E) \int_{E_{g+1}}^E \frac{\partial^2 \phi}{\partial E' \partial x}(\vec{r}, \vec{\Omega}, E') dE' \right\|^2 \\ &\leq \int_D \int_F \int_B \sum_{g=1}^G \chi_g(E) \left(\int_{E_{g+1}}^E \left| \frac{\partial^2 \phi}{\partial E' \partial x} \right| dE' \right)^2 dE d\vec{\Omega} d\vec{r} \end{aligned}$$

$$\begin{aligned}
& \leq \int_D \int_F \int_B \sum_{g=1}^G \chi_g(E) (E - E_{g+1}) \int_{E_{g+1}}^{E_g} \left| \frac{\partial^2 \phi}{\partial E' \partial x} \right|^2 dE' dE d\vec{\Omega} d\vec{r} \\
& = \int_D \int_F \sum_{g=1}^G \chi_g(E) \int_{E_{g+1}}^{E_g} \left| \frac{\partial^2 \phi}{\partial E' \partial x} \right|^2 dE' \frac{(E_g - E_{g+1})^2}{2} d\vec{\Omega} d\vec{r} \\
& \leq \max_{1 \leq g \leq G} \frac{(E_g - E_{g+1})^2}{2} \left\| \frac{\partial^2 \phi}{\partial E \partial x} \right\|^2 .
\end{aligned}$$

Applying the same procedure to each norm on the right hand side of equation (18),

$$\|\phi - z\|_{1,1,0}^2 \leq \max_{1 \leq g \leq G} \frac{(E_g - E_{g+1})^2}{2} \left[\left\| \frac{\partial \phi}{\partial E} \right\|^2 + \sum_{i=1}^N \left\| \frac{\partial \phi}{\partial x_i} \right\|^2 \right] .$$

Hence,

$$\|\phi - z\|_{1,1,0} \leq \max_{1 \leq g \leq G} \frac{E_g - E_{g+1}}{\sqrt{2}} \left\| \frac{\partial \phi}{\partial E} \right\|_1 . \quad /$$

b. Case 2: $M = 2$. For $M = 2$, it is reasonable to set $b_{1g}(E) = 1.0$, $1 \leq g \leq G$, as in case 1. The function $b_{2g}(E)$ may then be set equal to any continuous function of E on the interval (E_{g+1}, E_g) as long as $b_{2g}(E)$ is not identically equal to unity on that interval. If the flux, $\psi(\vec{r}, \vec{\Omega}, E)$, is expected to be a linear combination of two different spectra $b(E)$ and $h(E)$ then a good choice for $b_{2g}(E)$ may be $h(E)/b(E)$.

It is difficult to estimate $\|\phi - z\|_{1,1,0}$ for any arbitrary function $b_{2g}(E)$. For case 2 , $b_{2g}(E)$ will be set equal to E

which is a convenient choice for analysis purposes but as indicated above better choices may be available depending on the particular problem.

For this case, it is assumed that $\phi \in H^3(A)$. By assuming $\phi \in H^3(A)$, ϕ , any weak first-order derivative of ϕ , or any weak second-order derivative of ϕ is equivalent to a function which is absolutely continuous with respect to E for almost every $(\vec{r}, \vec{\Omega}) \in DXF$. The argument is analogous to that described in case 1.

For this case, z will be set equal to a linear interpolant in energy of ϕ over the energy intervals (E_{g+1}, E_g) for $g = 1, 2, \dots, G$. An estimate for $\|\phi - z\|_{1,1,0}$ is obtained in Theorem 3.

Theorem 3. Suppose $M = 2$, $\phi \in H^3(A)$, $b_{1g}(E) = 1$, $b_{2g}(E) = E$ for $1 \leq g \leq G$ and $\phi_{1g}(\vec{r}, \vec{\Omega})$, $\phi_{2g}(\vec{r}, \vec{\Omega})$ defined so that

$$z = \sum_{g=1}^G \chi_g(E) [\phi(\vec{r}, \vec{\Omega}, E_{g+1}) + \frac{E - E_{g+1}}{E_g - E_{g+1}} (\phi(\vec{r}, \vec{\Omega}, E_g) - \phi(\vec{r}, \vec{\Omega}, E_{g+1}))]$$

Then

$$\|\phi - z\|_{1,1,0} \leq \sqrt{\frac{7}{12}} \max_{1 \leq g \leq G} (E_g - E_{g+1})^2 \left\| \frac{\partial^2 \phi}{\partial E^2} \right\|_1.$$

Proof: Note that

$$\|\phi - z\|_{1,1,0} = \left\| \sum_{g=1}^G \chi_g(E) [\phi(\vec{r}, \vec{\Omega}, E) - \phi(\vec{r}, \vec{\Omega}, E_{g+1})] \right\| \quad (19)$$

$$\begin{aligned} & - \frac{E - E_{g+1}}{E_g - E_{g+1}} (\phi(\vec{r}, \vec{\Omega}, E_g) - \phi(\vec{r}, \vec{\Omega}, E_{g+1})) \|^2 \\ & + \sum_{i=1}^{N_x} \left\| \sum_{g=1}^G \chi_g(E) \left[\frac{\partial \phi}{\partial x_i}(\vec{r}, \vec{\Omega}, E) - \frac{\partial \phi}{\partial x_i}(\vec{r}, \vec{\Omega}, E_{g+1}) \right. \right. \\ & \left. \left. - \frac{\partial}{\partial x_i} \left\{ \frac{E - E_{g+1}}{E_g - E_{g+1}} (\phi(\vec{r}, \vec{\Omega}, E_g) - \phi(\vec{r}, \vec{\Omega}, E_{g+1})) \right\} \right] \right\|^2 \end{aligned}$$

Since $\phi \in H^3(A)$, ϕ or any first or second order weak derivative of ϕ is equivalent to a function which is absolutely continuous with respect to E . Thus, in addition to equations (16) and (17), the following equation is applicable:

$$\begin{aligned} \frac{\partial \phi}{\partial x_i}(\vec{r}, \vec{\Omega}, E) - \frac{\partial \phi}{\partial x_i}(\vec{r}, \vec{\Omega}, E_{g+1}) &= \\ &= \int_{E_{g+1}}^E \int_{E_{g+1}}^z \frac{\partial^3 \phi}{\partial E'^2 \partial x_i}(\vec{r}, \vec{\Omega}, E') dE' dz \\ &+ (E - E_{g+1}) \frac{\partial^2 \phi}{\partial E \partial x_i}(\vec{r}, \vec{\Omega}, E_{g+1}). \end{aligned}$$

Consider just one of the norms on the right hand side of equation (19). Any of them can be evaluated in a similar manner.

Consider

$$I^2 = \left\| \sum_{g=1}^G \chi_g(E) \left[\frac{\partial \phi}{\partial x} (\vec{r}, \vec{\Omega}, E) - \frac{\partial \phi}{\partial x} (\vec{r}, \vec{\Omega}, E_{g+1}) \right. \right. \\ \left. \left. - \frac{E - E_{g+1}}{E_g - E_{g+1}} \left(\frac{\partial \phi}{\partial x} (\vec{r}, \vec{\Omega}, E_g) - \frac{\partial \phi}{\partial x} (\vec{r}, \vec{\Omega}, E_{g+1}) \right) \right] \right\|^2.$$

Substituting into the above expression, it follows that

$$I^2 = \int_D \int_F \int_B \sum_{g=1}^G \chi_g(E) \left\{ \int_{E_{g+1}}^E \int_{E_{g+1}}^Z \frac{\partial^3 \phi}{\partial E'^2 \partial x} (\vec{r}, \vec{\Omega}, E') dE' dz \right. \\ \left. - \frac{E - E_{g+1}}{E_g - E_{g+1}} \int_{E_{g+1}}^{E_g} \int_{E_{g+1}}^Z \frac{\partial^3 \phi}{\partial E'^2 \partial x} (\vec{r}, \vec{\Omega}, E') dE' dz \right\}^2 dE d\vec{\Omega} d\vec{r}.$$

But since $(a - b)^2 \leq 2a^2 + 2b^2$,

$$I^2 \leq 2 \int_D \int_F \int_B \sum_{g=1}^G \chi_g(E) \left\{ \left(\int_{E_{g+1}}^E \int_{E_{g+1}}^Z \frac{\partial^3 \phi}{\partial E'^2 \partial x} dE' dz \right)^2 \right. \\ \left. + \left(\frac{E - E_{g+1}}{E_g - E_{g+1}} \right)^2 \left(\int_{E_{g+1}}^{E_g} \int_{E_{g+1}}^Z \frac{\partial^3 \phi}{\partial E'^2 \partial x} dE' dz \right)^2 \right\} dE d\vec{\Omega} d\vec{r}$$

By using the Schwarz inequality twice,

$$\left(\int_{E_{g+1}}^E \int_{E_{g+1}}^Z f(E') dE' dz \right)^2 \\ \leq (E - E_{g+1}) \int_{E_{g+1}}^E (z - E_{g+1}) \int_{E_{g+1}}^Z f^2(E') dE' dz$$

$$\leq \frac{(E - E_{g+1})^3}{2} \int_{E_{g+1}}^{E_g} f^2(E') dE' .$$

Hence,

$$\begin{aligned} I^2 &\leq 2 \int_D \int_F \int_B \sum_{g=1}^G \chi_g(E) \left[\frac{(E - E_{g+1})^3}{2} + \frac{(E - E_{g+1})^2}{2} (E_g - E_{g+1}) \right] \\ &\quad \cdot \int_{E_{g+1}}^{E_g} \left(\frac{\partial^3 \phi}{\partial E'^2 \partial x} \right)^2 dE' dE d\vec{\Omega} d\vec{r} \\ &\leq \frac{7}{12} \int_D \int_F \sum_{g=1}^G (E_g - E_{g+1})^4 \int_{E_{g+1}}^{E_g} \left(\frac{\partial^3 \phi}{\partial E'^2 \partial x} \right)^2 dE d\vec{\Omega} d\vec{r} \\ &\leq \frac{7}{12} \max_{1 \leq g \leq G} (E_g - E_{g+1})^4 \left\| \frac{\partial^3 \phi}{\partial E'^2 \partial x} \right\|^2 . \end{aligned}$$

Applying the same procedure to each norm on the right hand side of equation (19), the following inequality is obtained from which the desired result immediately follows.

$$\begin{aligned} \|\phi - z\|_{1,1,0}^2 &\leq \frac{7}{12} \max_{1 \leq g \leq G} (E_g - E_{g+1})^4 \cdot \\ &\quad \cdot \left[\left\| \frac{\partial^2 \phi}{\partial E^2} \right\|^2 + \sum_{i=1}^{N_x} \left\| \frac{\partial^3 \phi}{\partial E^2 \partial x_i} \right\|^2 \right] . \end{aligned}$$

B. Step 2 - Analysis of $\inf_{u^h \in S^h} \|z - u^h\|_{1,1,0} .$

Once the terms $b_{mg}(E)$ and $\phi_{mg}(\vec{r}, \vec{\Omega})$ are properly selected, the approximation by z through u^h reduces to an approximation in $H^1(\text{DXF})$.

By Minkowski's inequality,

$$\begin{aligned} \|z - u^h\|_{1,1,0} &\leq \sum_{m=1}^M \sum_{g=1}^G \|\chi_g(E) b_{mg}(E) \\ &\quad \cdot (\phi_{mg}(\vec{r}, \vec{\Omega}) - u_{mg}^h(\vec{r}, \vec{\Omega}))\|_{1,1,0} \\ &= \sum_{m=1}^M \sum_{g=1}^G c_{mg} \|\phi_{mg}(\vec{r}, \vec{\Omega}) - u_{mg}^h(\vec{r}, \vec{\Omega})\|_1 \end{aligned}$$

where

$$c_{mg} = \left(\int_{E_{g+1}}^{E_g} b_{mg}^2(E) dE \right)^{\frac{1}{2}}$$

and $\|f(\vec{r}, \vec{\Omega})\|_1$ represents the norm of $f(\vec{r}, \vec{\Omega})$ in $H^1(\text{DXF})$.

Thus,

$$\|\phi - \phi^h\| \leq c_0 \|\phi - z\|_{1,1,0} + \quad (20)$$

$$+ c_0 \sum_{m=1}^M \sum_{g=1}^G c_{mg} \inf_{u_{mg}^h \in S_1^h} \|\phi_{mg} - u_{mg}^h\|_1.$$

where

$$S_1^h = \{u_{mg}^h \in H^1(\text{DXF}): u_{mg}^h = \sum_{j=1}^J \xi_{jmg} \gamma_j(\vec{r}, \vec{\Omega}) \text{ and}$$

u_{mg}^h satisfies the boundary conditions on Γ_1 -

and on $\Gamma_2 \}$.

Estimates for $\|\phi - z\|_{1,1,0}$ have been given above. Estimates for $\inf_{u_{mg}^h \in S_1^h} \|\phi_{mg} - u_{mg}^h\|_1$ have been investigated for various geometries and finite elements by several authors. In addition, Sanchez and McCormick² summarize that the interpolation error generally varies as ρ_c^{k+1}/ρ_i where ρ_c is the radius of the smallest circle circumscribing the element and ρ_i is the radius of the largest circle inscribed within the element and k depends on the order of the polynomial approximation. Ciarlet²⁵ presents rigorous mathematical results supporting this statement provided that the function being approximated, ϕ_{mg} , is an element of $H^{k+1}(\text{DXF})$.

CHAPTER IV

MULTIGROUP EQUATIONS

Standard methods for solving the energy-dependent neutron transport equation first involve deriving a set of coupled equations, called the multigroup equations, by integrating the energy-dependent transport equation over energy intervals. The equations are coupled through the integral term in the transport equation which mathematically describes neutron transfers to one energy from other energies. Each multigroup equation is first solved numerically for the flux density with respect to the position and direction variables. The equations are then solved in an iterative manner over all energy groups.

Using a Galerkin approach, a set of multigroup equations that generalize the standard multigroup equations has not been derived before the present investigation. However, several authors have indicated that use of such an approach with proper basis functions does lead to alternative forms of multigroup equations.^{22,30}

To derive the standard multigroup equations using a Galerkin approach, a global approximation in the energy variable, $b(E)$, must first be introduced as was done in Chapter II. Using a Galerkin approach to approximate $\phi(\vec{r}, \vec{\Omega}, E)$, where $\psi(\vec{r}, \vec{\Omega}, E) = b(E) \phi(\vec{r}, \vec{\Omega}, E)$, results in a generalization of the standard multigroup equations.

In the first section of this chapter, a set of multigroup equations which generalize the standard multigroup equations are obtained. To derive these multigroup equations, the integrations in the inner

products in the Galerkin form of the transport equation are performed with respect to energy. Group cross sections and group flux densities are defined such that a set of Galerkin multigroup equations results from these integrations. If $M = 1$ and a step function approximation in the energy variable is used, then the standard multigroup equations are obtained. If $M = 2$ and a linear approximation in energy is used over each energy interval, a new form of multigroup equations is obtained. In Chapter VII, numerical examples are presented which indicate the increased accuracy that can be obtained using these multigroup equations. Also described in the first section of this chapter is an iterative procedure that can be used to solve these multigroup equations.

In the second section of this chapter, a new convergence result is obtained assuming that a Gauss-Seidel procedure is applied to iteratively solve the generalized set of multigroup equations.

1. Derivation

As defined earlier, the Galerkin equation has the following form with $\phi^h \in S^h$ an approximation to $\phi \in H^{1,1,0}(A)$.

$$(b\vec{\Omega} \cdot \nabla \phi^h, w) + (K\phi^h, w) = (q, w) \quad \text{for every } w \in S^h. \quad (21)$$

Let

$$\phi^h = \sum_{m=1}^M \sum_{g'=1}^G \chi_{g'}(E) b_{mg'}(E) \phi_{mg'}^h(\vec{r}, \vec{\Omega}) \quad (22)$$

$$\text{where } \phi_{mg'}^h(\vec{r}, \vec{\Omega}) = \sum_{j=1}^J \xi_{jmg'} \gamma_j(\vec{r}, \vec{\Omega}) .$$

As noted in Chapter III, $\gamma_j \chi_{g'} b_{mg'}$ for $1 \leq g' \leq G$, $1 \leq m \leq M$, $1 \leq j \leq J$ are linearly independent functions. Let $w = \chi_g b_{ng} \gamma_j$ in equation (21). Then equation (21) becomes:

$$(\vec{b}\vec{\Omega} \cdot \nabla \phi^h + K \phi^h, \chi_g b_{ng} \gamma_j) = (q, \chi_g b_{ng} \gamma_j) \quad (23)$$

$$\text{for } 1 \leq n \leq M, 1 \leq j \leq J, 1 \leq g \leq G .$$

The integrations over energy in the inner products in equation (23) are performed resulting in equation (24). The inner products in equation (24) therefore only contain integrations over $\vec{r}$ and $\vec{\Omega}$. Consequently, inner products involving only integrations with respect to the position and direction variables are used in the remainder of this section.

$$\begin{aligned} & \left(\sum_{m=1}^M [\vec{\Omega} \cdot \nabla \phi_{mg}^h(\vec{r}, \vec{\Omega}) a_{mng} + \sigma_{mng}(\vec{r}) a_{mng} \phi_{mg}^h(\vec{r}, \vec{\Omega}) \right. \\ & \quad \left. - \int_F \sum_{g'=1}^G \sigma f_{mng'g}(\vec{r}, \vec{\Omega}' \rightarrow \vec{\Omega}) c_{mg'} \phi_{mg'}^h(\vec{r}, \vec{\Omega}') d\vec{\Omega}' \right] , \\ & \quad \gamma_j(\vec{r}, \vec{\Omega}) \end{aligned} \quad (24)$$

$$= (q_{ng}(\vec{r}, \vec{\Omega}), \gamma_j(\vec{r}, \vec{\Omega}))$$

$$\text{for } 1 \leq n \leq M, 1 \leq g \leq G,$$

where

$$q_{ng}(\vec{r}, \vec{\Omega}) = \int_{E_{g+1}}^{E_g} q(\vec{r}, \vec{\Omega}, E) b_{ng}(E) dE$$

$$a_{mng} = \int_{E_{g+1}}^{E_g} b(E) b_{mg}(E) b_{ng}(E) dE$$

$$\sigma_{mng}(\vec{r}) = \frac{1}{a_{mng}} \int_{E_{g+1}}^{E_g} \sigma(\vec{r}, E) b(E) b_{ng}(E) b_{mg}(E) dE$$

$$c_{mg} = \int_{E_{g+1}}^{E_g} b(E) b_{mg}(E) dE$$

and

$$\begin{aligned} \sigma_{f_{mng},g}(\vec{r}, \vec{\Omega}' \rightarrow \vec{\Omega}) &= \frac{1}{c_{mg'}} \int_{E_{g+1}}^{E_g} b_{ng}(E) \\ &\cdot \int_{E_{g'+1}}^{E_{g'}} \sigma_f(\vec{r}; \vec{\Omega}', E' \rightarrow \vec{\Omega}, E) b(E') b_{mg'}(E') dE' dE. \end{aligned}$$

Equation (24) is the Galerkin form of the generalized multigroup neutron transport equation. Let $\psi^h(\vec{r}, \vec{\Omega}, E) = b(E) \phi^h(\vec{r}, \vec{\Omega}, E)$ be the approximation to $\psi(\vec{r}, \vec{\Omega}, E)$. Then,

$$\psi_g^h(\vec{r}, \vec{\Omega}) = \int_{E_{g+1}}^{E_g} b(E) \phi^h(\vec{r}, \vec{\Omega}, E) dE = \sum_{m=1}^M c_{mg} \phi_{mg}^h(\vec{r}, \vec{\Omega})$$

is the approximation to the group neutron flux density for energy interval (E_{g+1}, E_g) . Note that for $M = 1$, equation (24) reduces to the following familiar form:

$$\begin{aligned} & ([\vec{\Omega} \cdot \nabla \psi_g^h(\vec{r}, \vec{\Omega}) + \sigma_g(\vec{r}) \psi_g^h(\vec{r}, \vec{\Omega}) \\ & - \int_F \sum_{g'=1}^G \sigma_{f,g'g}(\vec{r}, \vec{\Omega}' \rightarrow \vec{\Omega}) \psi_{g'}^h(\vec{r}, \vec{\Omega}') d\vec{\Omega}' - q_{1g}(\vec{r}, \vec{\Omega})] , \\ & \gamma_j(\vec{r}, \vec{\Omega}) = 0 . \end{aligned} \quad (25)$$

Equation (25) is the standard Galerkin form of the multigroup equations assuming the separable approximation, i.e.

$$\psi^h(\vec{r}, \vec{\Omega}, E) = b(E) \sum_{g=1}^G \chi_g(E) \phi_{1g}^h(\vec{r}, \vec{\Omega}) .$$

However, there may be several advantages in having $M > 1$. For example, if the energy spectrum changes significantly through a region then the capability to accurately model the neutron flux density is improved if $M > 1$, especially if the spectrum is a linear combination of two or more different known spectra. Although the complexity of the calculations as well as storage requirements increases as M increases, the improved accuracy may offset these disadvantages.

By suitably defining vectors and matrices, the notation of equation (24) can be considerably simplified. Let A_g , B_g , and $C_{g'g}$ be $M \times M$ matrices where $1 \leq g, g' \leq G$. Define

$$[A_g]_{nm} = a_{mng}$$

$$[B_g]_{nm} = \sigma_{mng}(\vec{r}) a_{mng}$$

$$\text{and } [C_{g'g}]_{nm} = \sigma_{mng'g}(\vec{r}, \vec{\Omega}' \rightarrow \vec{\Omega}) c_{mg'}. .$$

Then equation (24) becomes:

$$(\vec{\Omega} \cdot \nabla A_g \phi_g + B_g \phi_g - \int_F \sum_{g'=1}^G C_{g'g} \phi_{g'} d\vec{\Omega}', \gamma_j) = (s_g, \gamma_j) \quad (26)$$

where ϕ_g and s_g are vectors of length M and

$$[\phi_g]_i = \phi_{ig}^h(\vec{r}, \vec{\Omega}) \quad \text{and} \quad [s_g]_i = q_{ig}(\vec{r}, \vec{\Omega}).$$

It can readily be shown that A_g , for $1 \leq g \leq G$, is positive definite and hence invertible. Therefore, equation (27) follows from equation (26) by multiplying both sides of equation (26) by the inverse of A_g .

$$(\vec{\Omega} \cdot \nabla \phi_g + \sigma_g \phi_g - \int_F \sum_{g'=1}^G \sigma_{g'g} \phi_{g'} d\vec{\Omega}', \gamma_j) = (q_g, \gamma_j) \quad (27)$$

$$\text{for } 1 \leq g \leq G, 1 \leq j \leq J$$

where

$$\sigma_g = A_g^{-1} B_g$$

$$\sigma_{g'g} = A_g^{-1} C_{g'g}$$

$$\text{and } q_g = A_g^{-1} s_g .$$

Note that $\sigma_g(\vec{r}, \vec{\Omega})$ and $\sigma_{g'g}(\vec{r}, \vec{\Omega}' \rightarrow \vec{\Omega})$ are $M \times M$ matrices and $\phi_g(\vec{r}, \vec{\Omega})$ and $q_g(\vec{r}, \vec{\Omega})$ are vectors of length M .

Equation (27) is of the same form as equation (25) except that in equation (27) the cross sections are arranged in matrices and the flux densities and sources are elements of vectors. Equation (27) represents a generalized form of the standard multigroup equation. The standard multigroup equations are obtained by letting $M = 1$ in equation (27).

The equations are rearranged for an iterative solution of the system. Let $L_{gg} \phi_g = \vec{\Omega} \cdot \nabla \phi_g + \sigma_g \phi_g - \int_F \sigma_{gg} \phi_g d\vec{\Omega}'$ and $L_{gg'} \phi_{g'} = - \int_F \sigma_{g'g} \phi_{g'} d\Omega'$ for $g' \neq g$. Then equation (27) is equivalent to:

$$\sum_{g'=1}^G (L_{gg'} \phi_{g'}, \gamma_j) = (q_g, \gamma_j) \quad \text{for } 1 \leq j \leq J, 1 \leq g \leq G .$$

Simplifying the notation further, $(L\vec{\phi}, \gamma_j) = (\vec{q}, \gamma_j)$ where

$[\vec{\phi}]_g = \phi_g$, $[\vec{q}]_g = q_g$, and $[L]_{gg'} = L_{gg'}$. Let $L = D + T_L + T_U$ where D is a diagonal matrix, T_L is lower triangular, and T_U is upper triangular. Then,

$$(D\vec{\phi}, \gamma_j) = -(T_L \vec{\phi}, \gamma_j) - (T_U \vec{\phi}, \gamma_j) + (\vec{q}, \gamma_j) \quad \text{for } 1 \leq j \leq J.$$

Now suppose an initial guess for $\vec{\phi}$, say $\vec{\phi}^1$, and consider the following iterative procedure:

$$(D\vec{\phi}^{n+1}, \gamma_j) = -(T_L \vec{\phi}^{n+1}, \gamma_j) - (T_U \vec{\phi}^n, \gamma_j) + (\vec{q}, \gamma_j) \quad (28)$$

$$\text{for } 1 \leq j \leq J \quad \text{and } n = 1, 2, 3, \dots$$

Equation (28) can be solved for $\vec{\phi}^{n+1}$ for $n \geq 1$ until specified convergence criteria are satisfied.

Convergence of the iterative procedure defined by equation (28) is considered in the next section. Note that the functions $\gamma_j(\vec{r}, \vec{\Omega})$, $1 \leq j \leq J$, have not been defined explicitly in these equations. These functions will be specified in Chapters V and VI for plane, spherical, and cylindrical one-dimensional geometry.

Before proceeding further, three remarks are made concerning the equations derived in this section.

First, an equation similar in form to equation (28) can be derived for the criticality problem which corresponds to a generalized eigenvalue problem. The iterative method then corresponds to a form of the power method and the largest eigenvalue along with the corresponding eigenvector is the converged solution.

Second, even- and odd-parity equations can be formulated for equation (27) as in reference (1) for the one-group transport equation except that in equation (27), ϕ_g is a vector of length M . This would be interesting to pursue since variational principles may be applicable to the even-parity form of equation (27).

Third, based on the multigroup transport equation given by equation (27), a form of multigroup P_1 equations as well as a form of multigroup diffusion equations can be derived. An approach can be used in this derivation that is similar to that given in reference (23) for deriving the standard P_1 and diffusion equations. Of course for $M = 1$, the derived equations reduce to the standard equations.

2. Convergence Result

The iterative procedure described by equation (28) is a Gauss-Seidel iterative procedure. Such iterative procedures are commonly used in practice to solve standard multigroup transport equations such as multigroup discrete ordinates equations. Each iteration described by equation (28) is referred to as an outer iteration. Additionally, an inner iteration is the solution for flux densities in each energy group over space and direction variables, if determined iteratively. A combination of inner and outer iterations is the standard procedure for solving multigroup transport equations.

In practice, a variety of techniques are available for accelerating convergence of inner and outer iterations.^{2,6,14,31,23} Examples are the rebalance, Chebyshev, and synthetic methods. Rebalance methods are based on enforcing neutron conservation (i.e. making neutron sources equal neutron losses) by multiplying the calculated flux densities by appropriate factors after each iteration. The Chebyshev method is used to accelerate convergence of the Jacobi iteration method and is described in reference (6) where it is

applied to the discrete ordinates equations. Synthetic methods use lower order approximations to accelerate convergence of a higher order method. For example, a diffusion approximation may be used to accelerate convergence of a discrete ordinates calculation.

By appropriate selection of the functions, $\gamma_j(\vec{r}, \vec{\Omega})$, in equation (28), the resulting equations are similar to the finite-difference discrete ordinates equations. This is demonstrated in Chapters V and VI for one-dimensional geometries. Presently, the finite-difference discrete ordinates equations are solved iteratively over spatial intervals and directions. However, based on the results of the present investigation, it appears that the discrete ordinates equations can be solved in an efficient noniterative manner over intervals and directions. In other words, inner iterations may be unnecessary for certain multigroup discrete ordinates methods. This will be demonstrated for plane geometry in Chapter V. However, it may still be advantageous to employ outer iterations to reduce computer storage requirements.

For the remainder of this section convergence of the outer iteration using a Gauss-Seidel method is analyzed. The convergence result is stated in Theorem 4. It is likely that the convergence could be improved by an SOR technique³³ but convergence criteria are difficult to obtain and are not discussed here. In Theorem 4, the vector $\vec{\phi}^n$ has length MG with elements of the form $\phi_{mg}^n(\vec{r}, \vec{\Omega})$ where

$$\phi_{mg}^n(\vec{r}, \vec{\Omega}) = \sum_{j=1}^J \xi_{jmg}^n \gamma_j(\vec{r}, \vec{\Omega}) .$$

The energy-dependent function $\phi^n(\vec{r}, \vec{\Omega}, E)$ is determined from $\vec{\phi}^n$ and is defined to be:

$$\phi^n(\vec{r}, \vec{\Omega}, E) = \sum_{m=1}^M \sum_{g=1}^G \chi_g(E) b_{mg}(E) \phi_{mg}^n(\vec{r}, \vec{\Omega}) .$$

Theorem 4. Suppose $B = \{E \in \mathbb{R}^1 : 0 \leq E_{\min} \leq E \leq E_{\max} < \infty\}$ and equation (10) holds with $\gamma > 0$. If $\gamma > 8\pi(E_{\max} - E_{\min})c$ where c is the maximum value of $b(E')$ of $f(\vec{r}; \vec{\Omega}', E' \rightarrow \vec{\Omega}, E)$ for $\vec{\Omega}, \vec{\Omega}' \in F$, $r \in D$, $E_{\min} \leq E' \leq E_{g+1} \leq E \leq E_g$, and $g = 1, 2, \dots, G$, then ϕ^n converges to ϕ^h in the $L^2(A)$ -norm using the Gauss-Seidel procedure described by equation (28). This condition is sufficient but not necessary.

Note that the magnitude of c is proportional to the probability of a neutron gaining energy in a nuclear collision. Theorem 4 therefore indicates that if group to group upscattering cross sections are sufficiently small, then the iterative procedure converges. By proper selection of energy group widths and functions, $b_{mg}(E)$, it is likely that the group upscattering cross sections can be made small while maintaining calculational accuracy.

Proof: Let A be the matrix defined by

$$[A]_{g',g} = \begin{cases} A_g & \text{if } g' = g \\ 0 & \text{otherwise} \end{cases} \quad (29)$$

for $1 \leq g, g' \leq G$, where A_g is the $M \times M$ matrix defined in the first section of this chapter. Thus, A is a $G \times G$ matrix whose elements are $M \times M$ matrices. Rearranging equation (28) and multiplying by matrix A results in:

$$(AL \vec{\phi}^{n+1}, \gamma_j) = (AT_u \vec{\phi}^{n+1}, \gamma_j) - (AT_u \vec{\phi}^n, \gamma_j) + (A\vec{q}, \gamma_j) \quad (30)$$

$$\text{where } L = D + T_u + T_L.$$

As before let $\vec{\phi}$ be the solution of equation (28) or (30) if the iterative procedure converges where $\vec{\phi}$ is a vector of length MG with elements of the form $\phi_{ig}^h(\vec{r}, \vec{\Omega})$. Hence,

$$(AL\vec{\phi}, \gamma_j) = (A\vec{q}, \gamma_j). \quad (31)$$

Subtracting equations (30) and (31), the following result is obtained:

$$(AL \vec{e}^{n+1}, \gamma_j) = (AT_u (\vec{e}^{n+1} - \vec{e}^n), \gamma_j) \quad (32)$$

$$\text{where } \vec{e}^{n+1} = \vec{\phi} - \vec{\phi}^{n+1}.$$

Notice that each inner product expression in equation (32) is a vector of length MG . According to equations (23) to (28), an element of vector $(AL \vec{e}^{n+1}, \gamma_j)$ has the form:

$$(b\vec{\Omega} \cdot \nabla e^{n+1} + Ke^{n+1}, \chi_g b_{lg} \gamma_j)$$

where $e^{n+1}(\vec{r}, \vec{\Omega}, E) = \phi^n(\vec{r}, \vec{\Omega}, E) - \phi^{n+1}(\vec{r}, \vec{\Omega}, E)$ and the inner product includes integration over energy. The corresponding element of the vector $(AT_u(\vec{e}^{n+1} - \vec{e}^n), \gamma_j)$ has the form:

$$\left(\int_F \int_{E_{\min}}^{E_{g+1}} \sigma f(\vec{r}; \vec{\Omega}', E' \rightarrow \vec{\Omega}, E) b(E') (e^{n+1}(\vec{r}, \vec{\Omega}', E') - e^n(\vec{r}, \vec{\Omega}', E')) d\vec{\Omega}' dE', \chi_g b_{lg} \gamma_j \right).$$

Therefore, multiplying both sides of equation (32) by the transpose of the appropriate vector and summing both sides over j for $j = 1, 2, \dots, J$ results in the following expression:

$$(b\vec{\Omega} \cdot \nabla e^{n+1} + Ke^{n+1}, e^{n+1}) = \sum_{g=1}^G \left(\int_F \int_{E_{\min}}^{E_{g+1}} \sigma f b (e^{n+1} - e^n) d\vec{\Omega}' dE', \sum_{\ell=1}^M \sum_{j=1}^J \beta_{j\ell g}^{n+1} b_{lg} \chi_g \gamma_j \right) \quad (33)$$

where

$$e^{n+1}(\vec{r}, \vec{\Omega}, E) = \sum_{g=1}^G \sum_{\ell=1}^M \sum_{j=1}^J \beta_{j\ell g}^{n+1} b_{lg}(E) \chi_g(E) \gamma_j(\vec{r}, \vec{\Omega})$$

and $\beta_{j\lambda g}^{n+1}$ are appropriately defined constants.

By equation (10), the left hand side of equation (33) is greater than or equal to $\gamma \|e^{n+1}\|^2$. The right hand side of equation (33) is less than or equal to $c \left(\int_F \int_B |e^{n+1} - e^n| d\vec{\Omega}' dE' , |e^{n+1}| \right)$ where c is defined in the statement of the theorem. Thus, using the Schwarz inequality twice,

$$\begin{aligned} \gamma \|e^{n+1}\|^2 &\leq c \left\| \int_F \int_B |e^{n+1} - e^n| d\vec{\Omega}' dE' \right\| \|e^{n+1}\| \\ &\leq c_3 \|e^{n+1} - e^n\| \|e^{n+1}\| \end{aligned}$$

$$\text{where } c_3 = 4\pi(E_{\max} - E_{\min})c .$$

Hence, by the triangle inequality, $\gamma \|e^{n+1}\| \leq c_3 \|e^{n+1}\| + c_3 \|e^n\|$.

Therefore, $\|e^{n+1}\| \leq \frac{c_3}{\gamma - c_3} \|e^n\|$ and it follows that

$$\|\phi^h - \phi^{n+1}\| \leq \left(\frac{c_3}{\gamma - c_3} \right)^n \|\phi^h - \phi^1\| .$$

Thus, if $\gamma > 2c_3 = 8\pi(E_{\max} - E_{\min})c$, then ϕ^n converges to ϕ^h in the $L^2(A)$ -norm.

/

CHAPTER V

PLANE GEOMETRY

The first aim in this chapter is to show that Galerkin's method applied to the multigroup transport equation leads to generalizations of the widely-used finite-difference discrete ordinates equations for plane geometry. The second aim is to show that these equations can be efficiently solved noniteratively.

The finite-difference discrete ordinates equations have not, before the present investigation, been derived using a Galerkin method applied to both direction and position variables. However, it has been shown that the finite-difference discrete ordinates equations in plane geometry can be obtained using a Galerkin approach in the position variable after assigning a set of discrete directions in the direction variable.²⁰ The above approach is much less illuminating than the approach used in the present investigation.

In the method presented here, the neutron flux density is approximated by a continuous piecewise polynomial of any specified degree in the spatial variable and a polynomial of any specified degree in the direction variable. For a continuous piecewise linear approximation in the spatial variable, this method reduces to the standard finite-difference discrete ordinates equations. For a parabolic approximation over spatial intervals, this method appears to reduce to a fourth-order finite difference method that is described in reference (13). Higher-order approximations can readily be formulated using the equations derived in this chapter.

Because these generalized equations are derived directly from applying Galerkin's method to both position and direction variables, an interesting new result is obtained. Specifically, it appears optimal in a manner described in the first section of this chapter to choose N and L such that $L \leq N \leq L + 1$ where N is the number of quadrature points in the direction variable and L is the degree of the polynomial expansion of the scattering cross section.

Although Galerkin's method is applied to the first-order multi-group transport equation in references (15), (16), (20) and (21), the analysis in the present investigation differs considerably in the basis functions and numerical integration rules selected. The method used in the present investigation for the position variable was adapted from that used in reference (34) for solving first-order ordinary differential equations. In this method, the expansion functions in the position variable are polynomials of any desired degree over each spatial interval. The flux density is known at one endpoint of the interval, from the calculation in the previous interval, and evaluated using Galerkin's method at the other endpoint of the interval. This procedure continues through all spatial intervals. Therefore, a continuous piecewise polynomial approximation in the spatial variable is obtained. The expansion functions in the direction variable are selected to be Legendre polynomials which form an orthogonal basis in the interval $[-1, 1]$. The nodal points at which the flux density is approximated are the quadrature points in the numerical integration rules. The numerical integration rules are selected so that the inner

products resulting in Galerkin's method are exactly evaluated. However, as mentioned in Chapter III, the boundary conditions are only required to be satisfied at the nodal points.

In the second section of this chapter, an efficient noniterative method for solving the system of equations derived in the first section is described, specifically for a linear approximation over spatial intervals. Presently, the finite-difference discrete ordinates equations are solved in an iterative manner. The noniterative approach may be useful, especially perhaps if vector processing computers are used to solve these equations.

Numerical comparisons between iterative and noniterative methods are described in Chapter VII as well as numerical comparisons between the parabolic and linear approximations with respect to the position variable.

1. Derivation of Galerkin Equations

In this section, a set of equations are derived that generalize the finite-difference discrete ordinates equations for one-dimensional plane geometry. These equations form a linear system for the unknown values of the flux density at specified positions and directions. The arrangement of these equations for a matrix formulation of this linear system is addressed in the second section of this chapter.

In one-dimensional plane geometry, the medium has a uniform composition at any fixed position x and the medium extends to infinity in the $\vec{y}$ and $\vec{z}$ directions. Refer to Figure 1.

Because of symmetry, the angular neutron flux, at any energy, depends only on variables x and μ where $\mu = \cos\theta = \vec{\Omega} \cdot \hat{e}_x$, $\vec{\Omega}$ is the unit vector in direction of neutron motion, and $\hat{e}_x$ is the unit vector in the $\vec{x}$ direction.

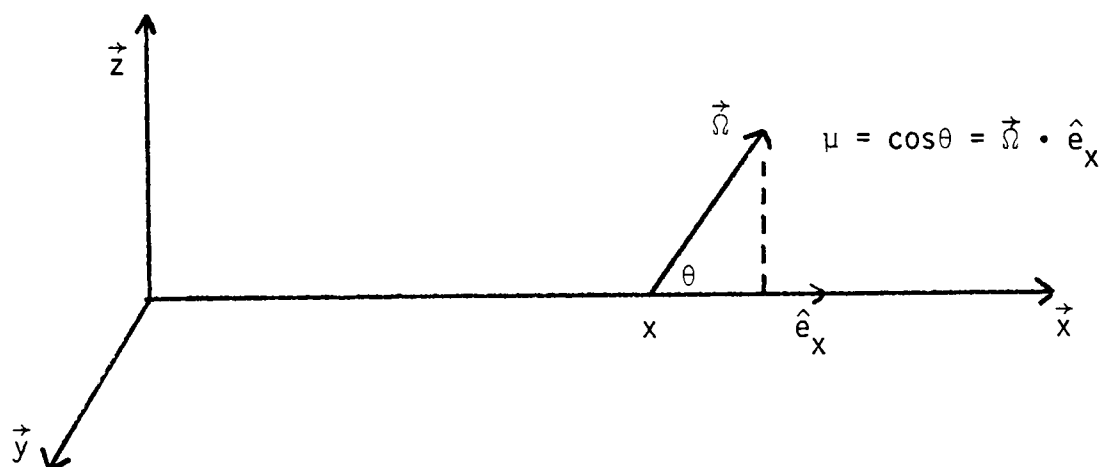


FIGURE 1. Illustration of One-Dimensional Plane Geometry

The Galerkin form of the multigroup transport equation is given by equation (27). For convenience, vector ϕ_g in equation (27) does not have a superscript h . For one-dimensional plane geometry, equation (27) has the following form:

$$\left(\mu \frac{\partial \phi_g}{\partial x}, \gamma_r\right) + (K_g \phi_g, \gamma_r) = (Q_g, \gamma_r) \quad (35)$$

$$\text{for } 1 \leq r \leq J, 1 \leq g \leq G,$$

where ϕ_g and Q_g are vectors of length M

$$[\phi_g]_m = \phi_{mg}^h(x, \mu) = \sum_{r=1}^J \varepsilon_{rmg} \gamma_r(x, \mu)$$

$$K_g \phi_g = \sigma_g(x) \phi_g(x, \mu) - \sum_{\ell=0}^{\infty} \frac{2\ell+1}{2} P_{\ell}(\mu) \sigma_{\ell g g}(x) \cdot$$

$$\cdot \int_{-1}^1 P_{\ell}(\mu') \phi_g(x, \mu') d\mu'$$

$$Q_g = q_g(x, \mu) + \sum_{\ell=0}^{\infty} \frac{2\ell+1}{2} P_{\ell}(\mu) \sum_{\substack{g'=1 \\ g' \neq g}}^G \sigma_{\ell g' g}(x) \cdot$$

$$\cdot \int_{-1}^1 P_{\ell}(\mu') \phi_{g'}(x, \mu') d\mu'$$

$$\sigma_{\ell g' g}(x) = 2\pi \int_{-1}^1 \sigma_{g' g}(x, \mu_0) P_{\ell}(\mu_0) d\mu_0$$

$P_{\ell}(\mu)$ is the ℓ th Legendre polynomial and σ_g and $\sigma_{\ell g' g}$ are $M \times M$ matrices.

Several remarks about equation (35) are necessary. The inner products in equation (35) are over x and μ where $a \leq x \leq c$ and $-1 \leq \mu \leq 1$. It is assumed in this section that the boundary at $x = a$ is reflective and the boundary at $x = c$ has a specified inward flux density. It is apparent from the equations presented how other boundary conditions can be treated. In equation (35), it is assumed that $\sigma_{g' g}(\vec{r}, \vec{\Omega}' \rightarrow \vec{\Omega})$ is a function of $\vec{r}$ and μ_0 where $\mu_0 = \vec{\Omega}' \cdot \vec{\Omega}$. This assumption is generally made in neutron transport problems. Equation (35) is obtained from equation (27) by expanding $\sigma_{g' g}(x, \mu_0)$

in Legendre polynomials, using the addition theorem for Legendre polynomials, and integrating over the azimuthal angle. A detailed description of this procedure is given in reference (23).

Note that $Q_g(x, \mu)$ includes the source into g due to scattering from other energy groups. Within-group scattering is included in the term $K_g \phi_g$.

Recall from the first section of Chapter IV that

$\psi_g^h(x, \mu) = \sum_{m=1}^M c_{mg} \phi_{mg}^h(x, \mu)$ is the approximation to the neutron flux density. Thus, $\psi_g^h(x, \mu)$ is the approximate neutron flux density of energy group g with direction μ at position x . If $M = 1$, equation (35) is the Galerkin form of the standard one-group neutron transport equation in plane geometry.

Since equation (35) may be solved over energy groups in the iterative manner indicated by equation (28), it remains to solve equation (35) over the position and direction variables for each energy group. Therefore, one energy group is considered and for notational convenience, the subscript g is dropped from the variables in equation (35). With this change, equation (35) has the following form:

$$\left(\mu \frac{\partial \phi}{\partial x} + K\phi, \gamma_r \right) = (Q, \gamma_r) \quad \text{for } 1 \leq r \leq J \quad (36)$$

$$\text{where } \phi(x, \mu) = \sum_{r=1}^J \beta_r \gamma_r(x, \mu)$$

and ϕ , Q , and β_r are vectors of length M .

The linearly independent functions, $\gamma_r(x, \mu)$, are selected to be piecewise polynomials in x , determined over partitions. These partitions are now introduced.

Define a partition Δ of $[a, c]$ to be the following:

$$\Delta = \{x_i \in [a, c], 0 \leq i \leq I : a = x_0 < x_1 < \dots < x_I = c\}.$$

Thus, there are I subintervals $[x_i, x_{i+1}]$ for $0 \leq i \leq I - 1$.

Each of these subintervals will be referred to as a major subinterval.

A set of I subpartitions of $[a, c]$ are defined based on the partition Δ . These subpartitions are defined as follows:

$$\Delta_i = \{x_{ip} \in [x_i, x_{i+1}], 0 \leq p \leq K:$$

$$x_i = x_{i0} < x_{i1} < \dots < x_{iK} = x_{i+1}\}$$

$$\text{for } 0 \leq i \leq I - 1.$$

Thus, each major subinterval is divided into K minor subintervals.

Therefore, altogether on $[a, c]$ there are KI minor subintervals and I major subintervals. It is convenient to choose the partition Δ such that the cross sections are independent of position on each major subinterval $[x_i, x_{i+1}]$, for $0 \leq i \leq I - 1$.

The functions $\gamma_r(x, \mu)$ are now defined. Let $\gamma_{ijn}(x, \mu) = P_n(\mu) \psi_{ji}(x) \chi_i(x)$ for $0 \leq i \leq I - 1$, $0 \leq n \leq N - 1$, and $1 \leq j \leq K$. Thus, $1 \leq r \leq INK$. In the above equation, $P_n(\mu)$ is the n th Legendre polynomial, $\psi_{ji}(x) = (x - x_i)^{j-1}$, and

$$\chi_i(x) = \begin{cases} 1 & \text{if } x_i \leq x \leq x_{i+1} \\ 0 & \text{otherwise.} \end{cases}$$

The functions $\psi_{ji}(x)$ should not be confused with the flux density $\psi(\vec{r}, \vec{\Omega}, E)$.

Define $\phi_i(x, \mu) = \sum_{n=0}^{N-1} \sum_{j=1}^{K+1} \beta_{ijn} P_n(\mu) \psi_{ji}(x) \chi_i(x)$ for $0 \leq i \leq I - 1$. Therefore, $\phi(x, \mu) = \sum_{i=0}^{I-1} \phi_i(x, \mu)$. The function

$\phi_i(x, \mu)$ is only nonzero over one major subdivision $[x_i, x_{i+1}]$.

The flux approximation is assumed to be continuous. Thus at the endpoints of each major subinterval,

$$\phi_i(x_i, \mu) = \phi_{i-1}(x_i, \mu) \quad (37)$$

Therefore, this approach produces a continuous piecewise polynomial approximation to the neutron flux density.

With the above selection for $\gamma_r(x, \mu)$, equation (36) has the form:

$$\left(\mu \frac{\partial \phi}{\partial x} + K\phi, P_n(\mu) \psi_{ji}(x) \chi_i(x)\right) = (Q, \chi_i(x) P_n(\mu) \psi_{ji}(x)) \quad (38)$$

$$\text{for } 0 \leq i \leq I - 1, 0 \leq n \leq N - 1, 1 \leq j \leq K.$$

Numerical integration rules are used to evaluate the inner products in equation (38). These rules are selected so that the inner products are exactly evaluated. It is assumed in this evaluation that Q has the form:

$$Q(x, \mu) = \sum_{i=0}^{I-1} Q_i(x, \mu)$$

$$\text{where } Q_i(x, \mu) = \sum_{n=0}^{N-1} \sum_{j=1}^{K+1} P_n(\mu) \psi_{ji}(x) \chi_i(x) s_{nji}.$$

It is shown that if an N -point Gaussian quadrature rule is used over the direction variable, μ , and a $K + 1$ -point Lobatto integration rule³⁵ is used over the position variable, x , in each major sub-interval, then the inner products are exactly evaluated. It is convenient to use the same integration rule to evaluate the integral that appears in the term $K\phi$. It is shown in the following proposition that the number of quadrature points, N , is related to the degree of the scattering expansion, L . (Note that in practice a finite number of terms are used in the polynomial expansion of the scattering cross section. In this paper, the integer L refers to the degree of the polynomial expansion of the scattering cross section.)

Proposition 3. If $L \leq N \leq L + 1$ where L is the degree of the polynomial expansion of the scattering cross section and N is the number of points in the Gaussian quadrature rule, then

$$K\phi = \sigma(x) \phi(x, \mu) - \sum_{\ell=0}^L \frac{2\ell+1}{2} P_{\ell}(\mu) \sigma_{\ell}(x) \cdot \quad (39)$$

$$\cdot \sum_{s=1}^N w_s \phi(x, \mu_s) P_{\ell}(\mu_s) .$$

Proof: Recall that $\phi(x, \mu)$ is a polynomial of degree $N - 1$ in μ . Also note that if $L \geq N - 1$, then

$$K\phi = \sigma(x) \phi(x, \mu) \quad (40)$$

$$- \sum_{\ell=0}^L \frac{2\ell+1}{2} P_{\ell}(\mu) \sigma_{\ell}(x) \int_{-1}^1 P_{\ell}(\mu') \phi(x, \mu') d\mu'$$

since $\int_{-1}^1 P_{\ell}(\mu') \phi(x, \mu') d\mu' = 0$ if $L > N - 1$ due to the orthogonality of the Legendre polynomials. In addition, if $N \geq L$ then

$$\int_{-1}^1 P_{\ell}(\mu') \phi(x, \mu') d\mu' = \sum_{s=1}^N w_s \phi(x, \mu_s) P_{\ell}(\mu_s)$$

where w_s and μ_s , for $1 \leq s \leq N$, are the Gaussian quadrature weights and points on $[-1, 1]$. This identity follows from the fact that N -point Gaussian quadrature is exact for polynomials of degree less than or equal to $2N - 1$.

Hence, if $L \leq N \leq L + 1$ then $K\phi$ can be expressed by equation (39). /

Proposition 3 indicates that this method will be successful even when the scattering is highly anisotropic, i.e. when the scattering cross section is strongly dependent on the scattering angle, since the exact contributions due to scattering from directions μ' to direction μ in equation (35) can be replaced by the equivalent form given by equation (39). (Of course, ϕ is assumed to be a polynomial of degree $N - 1$ in μ .)

Returning to equation (38), it is easy to see that the integration with respect to x is separated into I integrations, i.e. an integration over each major subinterval. Also, the first term in equation (38) can be integrated by parts with respect to x . Thus, equation (38) becomes:

$$\int_{-1}^1 \int_{x_i}^{x_{i+1}} L_{ji} \phi_i(x, \mu) P_n(\mu) dx d\mu = \int_{-1}^1 \int_{x_i}^{x_{i+1}} Q_i(x, \mu) \psi_{ji}(x) P_n(\mu) dx d\mu \quad (41)$$

$$\text{for } 0 \leq i \leq I - 1, 0 \leq n \leq N - 1, 1 \leq j \leq K,$$

where $L_{ji} \phi_i(x, \mu) =$

$$\begin{aligned} &= \frac{\mu}{x_{i+1} - x_i} (\phi_i(x_{i+1}, \mu) \psi_{ji}(x_{i+1}) - \phi_i(x_i, \mu) \psi_{ji}(x_i)) \\ &\quad - \mu \phi_i(x, \mu) \frac{d\psi_{ji}(x)}{dx} + K \phi_i(x, \mu) \psi_{ji}(x) \end{aligned}$$

and $K \phi_i$ is given by equation (39).

The integrands in equation (41) are polynomials of degree less than or equal to $2N - 1$ in μ and less than or equal to $2K - 1$ in x . Therefore, the integrals over x and μ can be evaluated exactly using a $K + 1$ -point Lobatto integration rule for integrals over x and an N -point Gaussian quadrature rule for the integrals over μ . Hence,

$$\begin{aligned} \sum_{r=1}^N w_r \sum_{p=0}^K L_{ji} \phi_i(x_{ip}, \mu_r) P_n(\mu_r) w_p &= \\ &= \sum_{p=0}^K \sum_{r=1}^N w_r Q_i(x_{ip}, \mu_r) w_p \psi_{ji}(x_{ip}) P_n(\mu_r) \end{aligned} \quad (42)$$

where μ_r and w_r for $1 \leq r \leq N$ are Gaussian quadrature points and weights and x_{ip} and w_p for $0 \leq p \leq K$ are Lobatto quadrature points and weights. Notice that the $K + 1$ points in each subpartition have been defined to be Lobatto quadrature points. A Lobatto quadrature rule is used so that two of the quadrature points are the endpoints of the major subinterval.

Equation (42) can be simplified considerably using Christoffel's identity.³⁶ Equation (43) is a consequence of this identity.

$$\sum_{n=0}^{N-1} P_n(\mu_r) P_n(\mu_m) \frac{2n+1}{2} = \begin{cases} \frac{1}{w_r} & \text{if } r = m \\ 0 & \text{otherwise.} \end{cases} \quad (43)$$

where μ_r, μ_m are quadrature points in the N -point Gaussian quadrature rule for the interval $[-1, 1]$ and w_r is the weight corresponding to point μ_r .

Multiplying both sides of equation (42) by $P_n(\mu_m) \frac{2n+1}{2}$, summing both sides of the equation over n for $0 \leq n \leq N-1$, and using equation (43) yields the following result:

$$\sum_{p=0}^K w_p L_{ji} \phi_i(x_{ip}, \mu_r) = \sum_{p=0}^K w_p Q_i(x_{ip}, \mu_r) \psi_{ji}(x_{ip}) \quad (44)$$

for $1 \leq r \leq N$, $1 \leq j \leq K$, and $0 \leq i \leq I-1$.

From equation (37) and assuming that the approximation satisfies the boundary conditions at the nodal points, the following equations are obtained:

$$\left\{ \begin{array}{ll} \phi_i(x_i, \mu_r) = \phi_{i-1}(x_i, \mu_r) & \text{for all } \mu_r \\ \phi_0(a, \mu_r) = \phi_0(a, -\mu_r) & \text{for all } \mu_r \\ \phi_I(c, \mu_r) = 0 & \text{for } \mu_r < 0. \end{array} \right. \quad (45)$$

Note that at $x = a$ a reflecting boundary condition is assumed and at $x = c$ an inward flux density is specified.

Equations (44) and (45) form a system of INK linear equations for determining the INK values of $\phi_i(x_{ip}, \mu_r)$. Also, recall that each ϕ_i and Q_i is a vector of length M where $M \geq 1$. Hence, the problem can be put in the form $A\phi = Q$ where A is a $INK \times INK$

matrix and each element of the vector ϕ or Q is a vector of length M . It turns out that the elements of ϕ and Q can be arranged such that A has bandwidth $(K + 1)N$. The matrix A is investigated in the next section specifically for $K = 1$.

Equations (44) and (45) reduce to the finite-difference discrete ordinates equations for $K = 1$ and $M = 1$ which is a linear approximation over mesh intervals. The case, $K = 1$, is considered in detail in the next section. Equations (44) and (45) reduce, for $K = 2$ and $M = 1$, to equations which appear to be the same as a new fourth-order finite difference method described in reference (13).

2. Noniterative Method of Solution

In this section, a noniterative method is described for solving equations (44) and (45) specifically for $K = 1$. However, a similar method is applicable for $K > 1$.

Setting $K = 1$ in equation (44) implies a two-point Lobatto integration rule. Therefore, $w_0 = w_1 = \frac{1}{2}$. Recalling that $\psi_{1i}(x) = 1$, equation (44) becomes

$$\left(\frac{\mu_r}{h_i} + \frac{\sigma_i}{2}\right) \phi_{i+1,r} + \left(\frac{-\mu_r}{h_i} + \frac{\sigma_i}{2}\right) \phi_{i,r} - \frac{1}{2} \sum_{s=1}^N w_s \sum_{\ell=0}^L \frac{2\ell+1}{2} \sigma_{\ell i} P_{\ell}(\mu_r) P_{\ell}(\mu_s) (\phi_{i+1,s} + \phi_{i,s}) = q_{i,r} \quad (46)$$

$$\text{for } 1 \leq r \leq N, 0 \leq i \leq I - 1$$

where for notational convenience:

$$h_i = x_{i+1} - x_i$$

$$q_{i,r} = \frac{1}{2}(Q_i(x_{i+1}, \mu_r) + Q_i(x_i, \mu_r))$$

$$\phi_{i,r} = \phi_i(x_i, \mu_r) = \phi_{i-1}(x_i, \mu_r)$$

$$\phi_{i+1,r} = \phi_i(x_{i+1}, \mu_r) = \phi_{i+1}(x_{i+1}, \mu_r)$$

$$\sigma_i = \text{value of } \sigma(x) \text{ in the } i\text{th interval}$$

$$\sigma_{li} = \text{value of } \sigma(x) \text{ in the } i\text{th interval.}$$

Note that the partition is chosen such that the cross sections are constant in each major subinterval. Also, note that the equations for $\phi_{i,r}$ and $\phi_{i+1,r}$ follow from equation (45) which specifies continuity of the flux density.

It is reasonable to solve equation (46) for $\phi_{i+1,r}$ if μ_r is positive and for $\phi_{i,r}$ if μ_r is negative. The resulting matrix will then have large diagonal elements. Using this approach, equations (47) and (48) follow directly from equation (46).

For $\mu_r > 0$ (47)

$$V_{ri} \phi_{i+1,r} + \sum_{\substack{s=1 \\ s \neq r}}^N Z_{sri} \phi_{i+1,s} + W_{ri} \phi_{i,r} + \\ + \sum_{\substack{s=1 \\ s \neq r}}^N Z_{sri} \phi_{i,s} = q_{i,r} .$$

For $\mu_r < 0$, (48)

$$W_{ri} \phi_{i,r} + \sum_{\substack{s=1 \\ s \neq r}}^N Z_{sri} \phi_{i,s} + V_{ri} \phi_{i+1,r} + \\ + \sum_{\substack{s=1 \\ s \neq r}}^N Z_{sri} \phi_{i+1,s} = q_{i,r} .$$

where

$$Z_{sri} = -w_s \sum_{\ell=0}^L \frac{2\ell+1}{4} \sigma_{\ell i} P_{\ell}(\mu_r) P_{\ell}(\mu_s)$$

$$V_{ri} = \frac{\mu_r}{h_i} + \frac{\sigma_i}{2} + Z_{rri}$$

and

$$W_{ri} = \frac{-\mu_r}{h_i} + \frac{\sigma_i}{2} + Z_{rri} \quad \text{for } 1 \leq r \leq N \text{ and } 0 \leq i \leq I-1.$$

Equations (47) and (48) can be put in the form $A\phi = q$.

For example, consider the case $N = 4$. Then let

$$\phi = \begin{bmatrix} \phi_{0,1} \\ \phi_{0,2} \\ \phi_{0,3} \\ \phi_{0,4} \\ \phi_{1,1} \\ \vdots \\ \vdots \\ \vdots \\ \phi_{I,2} \\ \phi_{I,3} \\ \phi_{I,4} \end{bmatrix} \quad \text{and} \quad q = \begin{bmatrix} 0 \\ 0 \\ q_{0,3} \\ q_{0,4} \\ q_{1,1} \\ \vdots \\ \vdots \\ \vdots \\ q_{I,2} \\ 0 \\ 0 \end{bmatrix}$$

where ϕ and q are vectors of length $N(I + 1)$ and the elements of ϕ and q are vectors of length M . Note that μ_1 and μ_2 are positive and μ_3 and μ_4 are negative.

Matrix A is given in Figure 2 for $N = 4$. In matrix A , the first two rows and last two rows each have one element which may be 0 or -1 depending on whether bare or reflecting boundary conditions are specified. In fact, matrix A could be shortened by four rows by incorporating the boundary conditions directly into the equations for $\phi_{0,3}$, $\phi_{0,4}$, $\phi_{1,1}$, $\phi_{1,2}$, $\phi_{I-1,3}$, $\phi_{I-1,4}$, $\phi_{I,1}$, and $\phi_{I,2}$. Also note that each element of matrix A is actually an $M \times M$ matrix.

$$A = \begin{bmatrix} 1 & 0 & 0 & (0 \text{ or } -1) & 0 & 0 & 0 & 0 & 0 & \dots \\ 0 & 1 & (0 \text{ or } -1) & 0 & 0 & 0 & 0 & 0 & 0 & \dots \\ Z_{130} & Z_{230} & W_{30} & Z_{430} & Z_{130} & Z_{230} & V_{30} & Z_{430} & 0 & \dots \\ Z_{140} & Z_{240} & Z_{340} & W_{40} & Z_{140} & Z_{240} & Z_{340} & V_{40} & 0 & \dots \\ W_{10} & Z_{210} & Z_{310} & Z_{410} & V_{10} & Z_{210} & Z_{310} & Z_{410} & 0 & \dots \\ Z_{120} & W_{20} & Z_{320} & Z_{420} & Z_{120} & V_{20} & Z_{320} & Z_{420} & 0 & \dots \\ 0 & 0 & 0 & 0 & Z_{131} & Z_{231} & W_{31} & Z_{431} & Z_{131} & \\ 0 & 0 & 0 & 0 & & & & & & \\ \cdot & \cdot & \cdot & \cdot & & & & & & \\ \cdot & \cdot & \cdot & \cdot & & & & & & \\ \cdot & \cdot & \cdot & \cdot & & & & & & \end{bmatrix}$$

FIGURE 2. Matrix A for $N = 4$.

Matrix A has the following equivalent form:

$$A = \begin{bmatrix} 1 & 0 & 0 & 0 & \dots & 0 \\ 0 & 1 & 0 & 0 & \dots & 0 \\ & A_0 & & & & \\ & & A_1 & & & \\ & & & A_2 & & \\ & & & & \ddots & \\ & & & & & A_{I-1} \\ & & & & & & \dots & 0 & 0 & 1 & 0 \\ & & & & & & \dots & 0 & 0 & 0 & 1 \end{bmatrix}$$

where each A_i , $0 \leq i \leq I - 1$, is a 4×8 matrix whose elements are $M \times M$ matrices. For the general case, each A_i is an $N \times 2N$ matrix whose elements are $M \times M$ matrices.

The properties of matrix A are investigated specifically for $M = 1$ but for an arbitrary value of N . It is shown that if h_i , $0 \leq i \leq I - 1$, are sufficiently small, and bare boundary conditions are specified, then A is likely to be strictly diagonally dominant. Hence, $A\phi = q$ can be solved using triangular decomposition without growth of roundoff errors.³⁷ Also, because of the banded structure of matrix A , storage requirements are low.

Consider equation (47) with $\mu_r > 0$. (A similar analysis can be performed using equation (48) with $\mu_r < 0$.) The diagonal element of A for $\mu_r > 0$ is:

$$\frac{\mu_r}{h_i} + \frac{\sigma_i}{2} - \frac{1}{2} w_r \sum_{\ell=0}^L \frac{2\ell+1}{2} \sigma_{\ell i} P_{\ell}^2(\mu_r) . \quad (49)$$

Let

$$\sigma'_{ir} = \sigma_i - w_r \sum_{\ell=0}^L \frac{2\ell+1}{2} \sigma_{\ell i} P_{\ell}^2(\mu_r) \quad (50)$$

and

$$T_{ri} = \sum_{\substack{s=1 \\ s \neq r}}^N |w_s \sum_{\ell=0}^L \frac{2\ell+1}{2} \sigma_{\ell i} P_{\ell}(\mu_r) P_{\ell}(\mu_s)| . \quad (51)$$

It follows from equation (47) that if

$$\left| \frac{\mu_r}{h_i} + \frac{\sigma'_{ir}}{2} \right| > \left| \frac{-\mu_r}{h_i} + \frac{\sigma'_{ir}}{2} \right| + T_{ri} \quad (52)$$

for every i and r , then matrix A is strictly diagonally dominant keeping in mind that a similar inequality is used if $\mu_r < 0$.

The following proposition provides conditions on h_i and σ'_{ir} to ensure that matrix A is strictly diagonally dominant.

Proposition 4. If $\sigma'_{ir} > T_{ri}$ and $\frac{2|\mu_r|}{h_i} > T_{ri}$ for $0 \leq i \leq I-1$

and $1 \leq r \leq N$, then matrix A is strictly diagonally dominant.

Proof: Assume in this proof that $\mu_r > 0$. A similar analysis can be made for $\mu_r < 0$.

Assume, for a given i and r , that $\frac{\mu_r}{h_i} \geq \frac{1}{2} \sigma'_{ir}$. Then if

$\sigma'_{ir} > T_{ri}$, equation (52) holds.

On the other hand assume, for a given i and r , that $\frac{1}{2} \sigma'_{ir} \geq \frac{\mu_r}{h_i}$. Then if $\frac{2\mu_r}{h_i} > T_{ri}$ equation (52) holds.

Summarizing, if $\frac{2\mu_r}{h_i} \geq \sigma'_{ir} > T_{ri}$ or $\sigma'_{ir} \geq \frac{2\mu_r}{h_i} > T_{ri}$ then equation (52) holds. Therefore, if σ'_{ir} and $\frac{2\mu_r}{h_i}$ are each greater than T_{ri} for every i and r , then matrix A is strictly diagonally dominant. /

First, note that h_i can be made sufficiently small so that $\frac{2|\mu_r|}{h_i} > T_{ri}$ for every i and r assuming that N is even to avoid any quadrature point, μ_r , from being equal to zero. Hence, by Proposition 4, it suffices to show that $\phi'_{ir} > T_{ri}$ for every i and r to prove that matrix A is strictly diagonally dominant (assuming that the h_i are sufficiently small).

The calculation required to show that $\sigma'_{ir} > T_{ri}$ for every i and r is tedious and therefore is not presented. The result of the calculation is that if the ratio of the absorption to scattering cross section for each spatial interval is sufficiently large then $\sigma'_{ir} > T_{ri}$ for every i and r . Hence, matrix A is strictly diagonally dominant provided that the absorption cross sections are sufficiently large and the interval widths, h_i , are sufficiently small.

However, the values of h_i required to satisfy the inequality in Proposition 4 may be too small for practical calculations requiring too much computer storage space. In addition, if the h_i are not

sufficiently small, the calculations may yield negative values for the flux densities. In practice, for the finite-difference discrete ordinates equations (i.e. $M = 1$ and $K = 1$), flux fix-up methods are used to recalculate flux densities if those flux densities are negative.

An alternative procedure is to use a higher order approximation, $K > 1$, in equation (44). These higher order methods are likely to be near positive, i.e. the appearance of negative fluxes is so infrequent for reasonable mesh sizes that the solution process is not significantly affected. This remark is based on the conclusion in reference (10) that the linear discontinuous method is near positive. (The linear discontinuous method is very similar to the method derived in the present investigation for $K = 2$ and according to reference (11) the linear discontinuous method also uses a continuous piecewise parabolic approximation with respect to the position variable.) The equations obtained in the present investigation for $K > 1$ are also adaptable to noniterative solution. Computational results using a noniterative procedure for $K = 2$ are described in Chapter VII.

Several other discrete ordinates methods, in addition to the method described in the present investigation, are adaptable to noniterative solution such as the method used in ONETRAN.¹⁴ In other words, noniterative procedures may be substituted for inner iterative procedures in these discrete ordinates methods. However, no research in this area is reported in the literature.

CHAPTER VI

SPHERICAL AND CYLINDRICAL GEOMETRY

With regard to transport approximations, the literature contains much less reported work for spherical and cylindrical geometries than for plane geometry. The reason for this is that in curved geometries, the angular coordinate of the flux density changes even in collision-free motion resulting in a more complicated form of the transport equation than for plane geometry. Also, the analysis of various transport approximations is complicated by auxiliary conditions that relate the flux densities between adjacent position and direction mesh points as in the diamond-difference discrete ordinates methods.

The aim in this chapter is to derive equations that relate flux densities at a set of specified directions and spatial positions for one-dimensional spherical and cylindrical geometries. In this derivation, a Galerkin approach is applied to the multigroup form of the transport equation in spherical and cylindrical geometries. No auxiliary conditions are assumed. The method used is similar to that used for plane geometry and the resulting equations are similar in form to equations (44) and (45) that were obtained for plane geometry. However, these equations do not simplify to the diamond-difference discrete ordinates equations for any order of approximation as was the case for plane geometry. However, as for plane geometry, these equations can be solved in an iterative or noniterative manner

over each energy group. The equations derived in this chapter have not been previously derived elsewhere.

The analysis presented in this chapter is somewhat briefer than that presented in the previous chapter for plane geometry since a noniterative method of solution is not specifically considered. In addition, numerical comparisons are not made. However, two interesting results are obtained that relate the number of discrete directions to the degree, L , of the scattering expansion. For spherical geometry, it appears optimal to select N such that $L \leq N \leq L + 1$ where N is the number of discrete directions. In cylindrical geometry, it appears optimal to choose $L \leq M \leq L + 1$ and $N \geq L$ where M and N are related to the total number of discrete directions and are defined in section 2. Also, in cylindrical geometry, the angular quadrature points are different from those generally used in discrete ordinates methods.

In this chapter, ϕ provides the one-group approximation to the neutron flux density and is only a function of position and direction. Therefore, the inner products in this chapter are only over position and direction. The one-group equation is coupled to the other groups in a multigroup calculation through the source term Q . In addition, ϕ can be a vector as for plane geometry.

1. Spherical Geometry

In spherical geometry, the analogous one-group equation to equation (35) has the form:

$$(r^2 \mu \frac{\partial \phi}{\partial r} + r(1 - \mu^2) \frac{\partial \phi}{\partial \mu} + r^2 K\phi, \gamma_\alpha) = (r^2 Q, \gamma_\alpha) \quad (53)$$

for $1 \leq \alpha \leq J$,

$$\text{where } K\phi = \sigma(r) \phi(r, \mu) - \sum_{\ell=0}^{\infty} \frac{2\ell+1}{2} P_\ell(\mu) \sigma_\ell(r) \\ \cdot \int_{-1}^1 P_\ell(\mu') \phi(r, \mu') d\mu'$$

and r = radius

where the subscript g has been dropped for convenience.

In one-dimensional spherical geometry, the medium is uniform in composition at any fixed distance r from a given point C which is the center of the sphere. See Figure 3. Because of symmetry, the angular neutron flux density, at any energy, depends only on variables r and μ where $\mu = \cos\theta = \vec{\Omega} \cdot \hat{e}_r$, $\vec{\Omega}$ is the unit vector in direction of neutron motion, and $\hat{e}_r$ is the unit vector in the $\vec{r}$ direction.

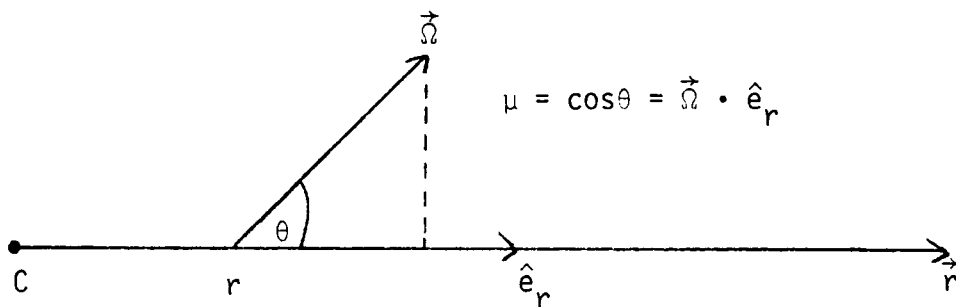


FIGURE 3. Illustration of One-Dimensional Spherical Geometry

Before continuing, note that for $M = 1$, equation (53) is the standard Galerkin form of the one-group neutron transport equation in spherical geometry in which both sides have been multiplied by r^2 . The factor r^2 originates from the volume element for spherical geometry.

In spherical geometry, $0 \leq r \leq c$ and $-1 \leq \mu \leq 1$. The boundary at $r = 0$ is reflective and the boundary at $r = c$ is assumed to have a specified inward flux density.

Similar to plane geometry, a partition Δ is defined as follows:

$$\Delta = \{r_i \in [0, c], 0 \leq i \leq I: 0 = r_0 < r_1 < \dots < r_I = c\}$$

where each $[r_i, r_{i+1}]$ for $0 \leq i \leq I - 1$ is defined to be a major subinterval. A collection of I subpartitions are defined as follows:

$$\Delta_i = \{r_{ip} \in [r_i, r_{i+1}], 0 \leq p \leq K:$$

$$r_i = r_{i0} < r_{i1} < \dots < r_{iK} = r_{i+1}\}$$

for $0 \leq i \leq I - 1$ where each subinterval in a subpartition is defined to be a minor subinterval. Therefore, on the interval $[0, c]$, there are IK minor subintervals and I major subintervals. The cross sections are assumed to be constant on each major subinterval.

In spherical geometry, the functions, $\gamma_\alpha(r, \mu)$, are similar to those for plane geometry. Let $\gamma_{nji}(r, \mu) = P_n(\mu) \psi_{ji}(r) \chi_i(r)$ for $0 \leq i \leq I - 1$, $0 \leq n \leq N - 1$, and $1 \leq j \leq K - 1$, where $P_n(\mu)$ is the n th Legendre polynomial, $\psi_{ji}(r) = (r - r_i)^{j-1}$, and

$$\chi_i(r) = \begin{cases} 1 & \text{if } r_i \leq r \leq r_{i+1} \\ 0 & \text{otherwise.} \end{cases}$$

The approximating function, $\phi(r, \mu)$, is defined as follows:

$$\phi(r, \mu) = \sum_{i=0}^{I-1} \phi_i(r, \mu) \quad (54)$$

$$\text{where } \phi_i(r, \mu) = \sum_{n=0}^{N-1} \sum_{j=1}^K P_n(\mu) \psi_{ji}(r) \chi_i(r) \xi_{nji}.$$

In addition, ϕ is continuous so that

$$\phi_i(r_i, \mu) = \phi_{i-1}(r_i, \mu) \quad (55)$$

Also, assume that $Q(r, \mu) = \sum_{i=0}^{I-1} Q_i(r, \mu)$ where

$$Q_i(r, \mu) = \sum_{n=0}^{N-1} \sum_{j=1}^K P_n(\mu) \psi_{ji}(r) \chi_i(r) s_{nji}.$$

Substituting the above expressions into equation (53) yields the following equation where the first two terms have been integrated by parts over the variables r and μ .

$$\int_{r_i}^{r_{i+1}} \int_{-1}^1 [L_{ji} \phi_i(r, \mu) P_n(\mu) - G_{ji} \phi_i(r, \mu) P'_n(\mu)] d\mu dr \quad (56)$$

$$= \int_{r_i}^{r_{i+1}} \int_{-1}^1 r^2 Q_i(r, \mu) P_n(\mu) \psi_{ji}(r) d\mu dr$$

for $0 \leq i \leq I - 1$, $1 \leq j \leq K - 1$, and $0 \leq n \leq N - 1$,

where

$$L_{ji} \phi_i(r, \mu) = -\phi_i(r, \mu) \mu r^2 \psi'_{ji}(r) + r^2 K\phi_i(r, \mu) \psi_{ji}(r)$$

$$+ \frac{1}{r_{i+1} - r_i} (r_{i+1}^2 \mu \phi_i(r_{i+1}, \mu) \psi_{ji}(r_{i+1})$$

$$- r_i^2 \mu \phi_i(r_i, \mu) \psi_{ji}(r_i))$$

$$G_{ji} \phi_i(r, \mu) = \phi_i(r, \mu) r \psi_{ji}(r) (1 - \mu^2)$$

$$\text{and } K\phi_i(r, \mu) = \sigma_i \phi_i(r, \mu) - \sum_{\ell=0}^L \frac{2\ell+1}{2} P_\ell(\mu) \sigma_{\ell i} \cdot$$

$$\cdot \sum_{s=1}^N w_s \phi(r, \mu_s) P_\ell(\mu_s) \cdot$$

Note in equation (56) that it is assumed that $L \leq N \leq L + 1$ since Proposition 3 has been used to simplify $K\phi_i$. Also, the cross sections in each major subinterval have been assumed to be constant.

The integrands in equation (56) are polynomials of degree $2K - 1$ in r and of degree $2N - 1$ in μ . By using a $K + 1$ -point Lobatto integration rule in r and an N -point Gaussian quadrature rule in μ , the integrals in equation (56) can be evaluated exactly. Note that a Lobatto integration rule is used in the position variable so

that two of the quadrature points are the endpoints of the major subinterval.

Performing these numerical integrations, equation (56) becomes:

$$\begin{aligned} & \sum_{t=1}^N w_t \sum_{p=0}^K w_p [L_{ji} \phi_i(r_{ip}, \mu_t) P_n(\mu_t) - G_{ji} \phi_i(r_{ip}, \mu_t) P'_n(\mu_t)] \quad (57) \\ & = \sum_{t=1}^N w_t \sum_{p=0}^K w_p r_{ip}^2 Q_i(r_{ip}, \mu_t) P_n(\mu_t) \psi_{ji}(r_{ip}) \end{aligned}$$

$$\text{for } 0 \leq i \leq I - 1, 1 \leq j \leq K - 1, 0 \leq n \leq N - 1,$$

where r_{ip} and w_p for $0 \leq p \leq K$ are Lobatto points and weights in the major subinterval $[r_i, r_{i+1}]$ and μ_t and w_t for $0 \leq t \leq N$ are Gaussian quadrature points and weights in the interval $[-1, 1]$.

Equation (57) can be simplified somewhat using equation (43). Multiplying both sides of equation (57) by $P_n(\mu_m) \frac{2n+1}{2}$ where μ_m is a Gaussian quadrature point, summing both sides of the equation over $0 \leq n \leq N - 1$, and using equation (43) results in the following equation:

$$\begin{aligned} & \sum_{p=0}^K w_p [L_{ji} \phi_i(r_{ip}, \mu_m) - \sum_{t=1}^N w_t G_{ji} \phi_i(r_{ip}, \mu_t) \cdot \\ & \quad \cdot \sum_{n=0}^{N-1} \frac{2n+1}{2} P_n(\mu_m) P'_n(\mu_t)] \quad (58) \\ & = \sum_{p=0}^K w_p r_{ip}^2 Q_i(r_{ip}, \mu_m) \psi_{ji}(r_{ip}) \end{aligned}$$

where, in addition

$$\left\{ \begin{array}{ll} \phi_i(r_i, \mu_m) = \phi_{i-1}(r_i, \mu_m) & \text{for all } \mu_m \\ \\ \phi_0(0, \mu_m) = \phi_0(0, -\mu_m) & \text{for all } \mu_m \\ \\ \phi_I(c, \mu_m) = 0 & \text{for } \mu_m < 0 \end{array} \right. \quad (59)$$

for $1 \leq m \leq N$, $1 \leq j \leq K - 1$, and $0 \leq i \leq I - 1$.

Notice that a reflecting boundary condition is specified at $r = 0$ and that an inward flux density is specified at $r = c$. However, the boundary conditions are only required to be satisfied at the nodal points, i.e. only at discrete directions on the boundary. Also note that w_p , $0 \leq p \leq K$, and w_t , $1 \leq t \leq N$, are Lobatto quadrature weights and Gaussian quadrature weights, respectively.

Equations (58) and (59) give $IN(K - 1)$ equations for the INK values of the flux density to be determined at the directions and positions specified by the numerical integration rules. The remaining IN additional equations can be found in the following manner.

For each value μ_m , $\phi_i(r, \mu_m)$ is a polynomial of degree $K - 1$ in r over the major subinterval $[r_i, r_{i+1}]$. However, in equation (58), $\phi_i(r, \mu_m)$ is required at $K + 1$ points in this subinterval. Therefore, $\phi_i(r, \mu_m)$ at one point, say r_{ic} , can be determined in terms of the values of $\phi_i(r, \mu_m)$ at the other K points in the major subinterval. This can be done for each of the I major subintervals

for each direction yielding IN additional equations. It appears reasonable to let c equal $\frac{1}{2}K$ if K is even and equal $\frac{1}{2}(K - 1)$ or $\frac{1}{2}(K + 1)$ if K is odd. Then the additional IN equations are (using Lagrange interpolation):

$$\phi_i(r_{ic}, \mu_m) = \sum_{\substack{p=0 \\ p \neq c}}^K \phi(r_{ip}, \mu_m) L_p(r_{ic}) \quad (60)$$

$$\text{for } 0 \leq i \leq I - 1, 1 \leq m \leq N$$

$$\text{where } L_p(r_{ic}) = \prod_{\substack{\ell=0 \\ \ell \neq p \\ \ell \neq c}}^K \left(\frac{r_{ic} - r_{i\ell}}{r_{ip} - r_{i\ell}} \right) .$$

Similarly,

$$Q_i(r_{ic}, \mu_m) = \sum_{\substack{p=0 \\ p \neq c}}^K Q_i(r_{ip}, \mu_m) L_p(r_{ic}) . \quad (61)$$

It appears that equations (58), (59), and (60) can be arranged such that the INK values of the flux density can be calculated by iterative or noniterative procedures. The equations seem complicated because the value of K is arbitrary. Actually for $K = 2$, which is a linear approximation over mesh intervals, the equations appear to be less complicated than the diamond-difference discrete ordinates equations which are also based on a linear approximation in the position variable. However, the above equations do not reduce to the diamond-difference equations for any value of K as was the case for

plane geometry. In addition, it appears that even for $K = 2$, equations (58), (59), and (60) may be more accurate than the diamond-difference discrete ordinates equations especially for values of N greater than two. The diamond-difference discrete ordinates equations contain an assumption of linearity between adjacent discrete directions. However, equation (58) is based on a polynomial approximation of degree $N - 1$ in μ . It appears that a numerical investigation is required to better compare the accuracy of this method with other discrete ordinates methods. Such a numerical investigation was not undertaken in the present work.

2. Cylindrical Geometry

In cylindrical geometry, the analogous one-group equation to equation (35) has the form:

$$\begin{aligned} & (r\sqrt{1-\mu^2} \cos\chi \frac{\partial\phi}{\partial r} - \sqrt{1-\mu^2} \sin\chi \frac{\partial\phi}{\partial\chi} + rK\phi, \gamma_\alpha) \\ & = (rQ, \gamma_\alpha) \quad \text{for } 1 \leq \alpha \leq J, \end{aligned} \quad (62)$$

where

$$\begin{aligned} \phi &= \phi(r, \mu, \chi) \quad \text{for } 0 \leq r \leq c \\ & \quad -1 \leq \mu \leq 1 \\ & \quad -\pi \leq \chi \leq \pi. \end{aligned}$$

In one-dimensional cylindrical geometry, the medium is uniform in composition at any fixed radial distance r from the z -axis (axis

of the cylinder). See Figure 4. Because of symmetry, the angular neutron flux density, at any energy, depends only on variables r , μ , and χ where $\mu = \vec{\Omega} \cdot \hat{e}_z$ and $(1 - \mu^2)^{\frac{1}{2}} \cos \chi = \vec{\Omega} \cdot \hat{e}_r$. The unit vectors $\hat{e}_z$ and $\hat{e}_r$ are in directions $\vec{z}$ and $\vec{r}$, respectively.

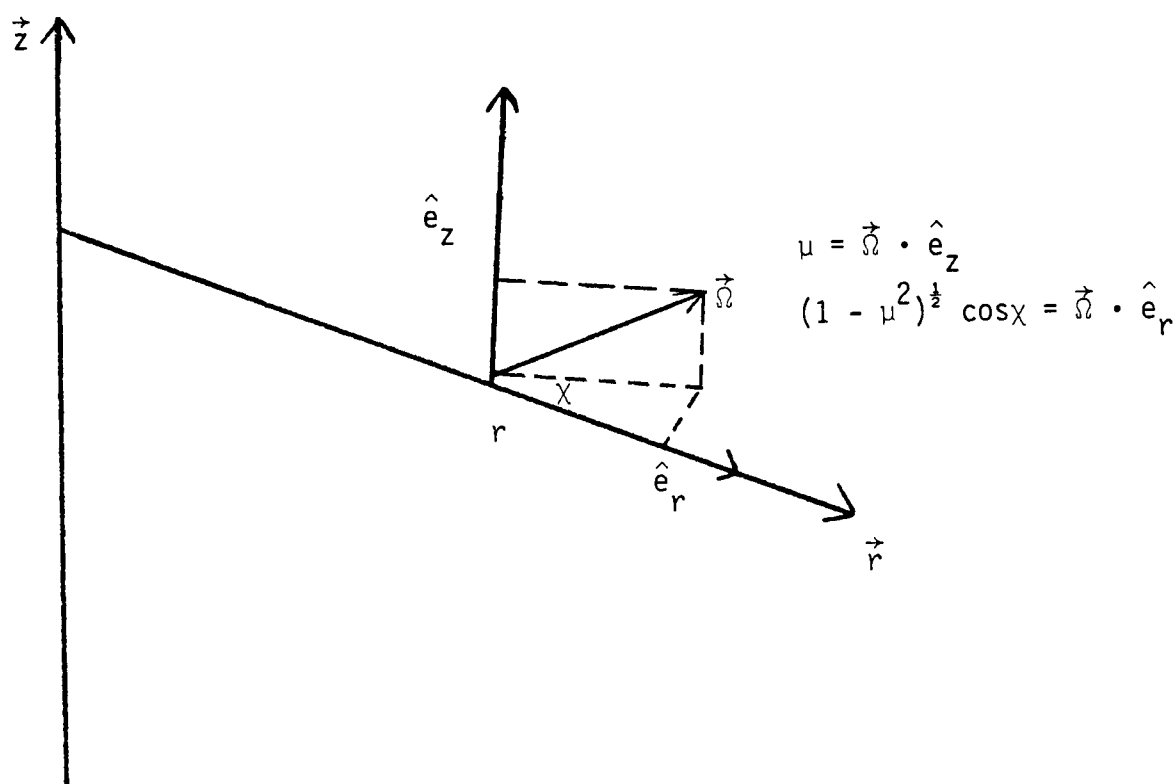


FIGURE 4. Illustration of One-Dimensional Cylindrical Geometry

By symmetry, $\phi(r, \mu, \chi) = \phi(r, -\mu, \chi) = \phi(r, \mu, -\chi) = \phi(r, -\mu, -\chi)$.

In addition, $\int_F d\vec{\Omega} = \int_{-1}^1 d\mu \int_{-\pi}^{\pi} d\chi = 4\pi$.

In cylindrical geometry, however, $K\phi$ is more complicated than in plane or spherical geometry as shown in the following equation:

$$K\phi(r, \mu, \chi) = \sigma(r) \phi(r, \mu, \chi) \quad (63)$$

$$\begin{aligned} & - \sum_{\ell=0}^{\infty} \frac{2\ell+1}{4\pi} \sigma_{\ell}(r) \int_{-1}^1 \int_{-\pi}^{\pi} \phi(r, \mu', \chi') \cdot \\ & \cdot [P_{\ell}(\mu) P_{\ell}(\mu') + 2 \sum_{v=1}^{\ell} \frac{(\ell-v)!}{(\ell+v)!} P_{\ell}^v(\mu) P_{\ell}^v(\mu') \cdot \\ & \cdot \cos v(\chi - \chi')] d\chi' d\mu' \end{aligned}$$

where $P_{\ell}^v(\mu)$ are the associated Legendre polynomials.

Before continuing, note that equation (62) is the standard Galerkin form of the one-group neutron transport equation in one-dimensional cylindrical geometry in which both sides of the equation have been multiplied by r . The factor r originates from the volume element for cylindrical geometry.

In cylindrical geometry, $0 \leq r \leq c$, $-1 \leq \mu \leq 1$, and $-\pi \leq \chi \leq \pi$. The boundary at $r = 0$ is reflective and the boundary at $r = c$ is assumed to have a specified inward flux density.

As for plane and spherical geometry, a partition Δ is defined as follows:

$$\Delta = \{r_i \in [0, c], 0 \leq i \leq I - 1:$$

$$0 = r_0 < r_1 < \dots < r_I = c\},$$

where each $[r_i, r_{i+1}]$ for $0 \leq i \leq I - 1$ is defined to be a major subinterval. A collection of I subpartitions are defined as follows:

$$\Delta_i = \{r_{ip} \in [r_i, r_{i+1}], 0 \leq p \leq K:$$

$$r_i = r_{i0} < r_{i1} < \dots < r_{iK} = c\}$$

for $0 \leq i \leq I - 1$ where each subinterval in a subpartition is defined to be a minor subinterval. Therefore, on the interval $[0, c]$, there are IK minor subintervals and I major subintervals. The cross sections are assumed to be constant on each major subinterval.

The functions, $\gamma_\alpha(r, \mu, \chi)$, are somewhat different for cylindrical geometry than for plane or spherical geometry. Let

$$\gamma_{njim}(r, \mu, \chi) = (1 - \mu^2)^{n/2} \cos m\chi \psi_{ji}(r) \chi_i(r) \text{ for } 0 \leq i \leq I - 1,$$

$0 \leq n \leq N - 1$, $0 \leq m \leq M - 1$, and $1 \leq j \leq K - 1$, where

$$\psi_{ji}(r) = (r - r_i)^{j-1} \text{ and}$$

$$\chi_i(r) = \begin{cases} 1 & \text{if } r_i \leq r \leq r_{i+1} \\ 0 & \text{otherwise.} \end{cases}$$

The function, $\phi(r, \mu, \chi)$, is defined as follows:

$$\begin{aligned} \phi(r, \mu, \chi) &= \sum_{i=0}^{I-1} \sum_{n=0}^{N-1} \sum_{j=1}^K \sum_{m=0}^{M-1} a_{njim} \gamma_{njim}(r, \mu, \chi) \quad (64) \\ &= \sum_{i=0}^{I-1} \phi_i(r, \mu, \chi) \end{aligned}$$

where ϕ is expanded in even functions of μ and χ since ϕ is even in μ and χ . Notice that the expression for ϕ includes

integral powers of the direction cosines in the $\vec{z}$ and $\vec{r}$ directions. In addition, ϕ is assumed to be continuous

$$\phi_i(r_i, \mu, \chi) = \phi_{i-1}(r_i, \mu, \chi) \quad (65)$$

Also, assume that $Q(r, \mu, \chi) = \sum_{i=0}^{I-1} Q_i(r, \mu, \chi)$ where

$$Q_i(r, \mu, \chi) = \sum_{n=0}^{N-1} \sum_{j=1}^{K-1} \sum_{m=0}^{M-1} s_{njim} \gamma_{njim}(r, \mu, \chi) .$$

Before deriving equations that are similar to equations (44) and (45) by evaluating equation (62) using numerical integration rules, numerical integration rules are selected that exactly evaluate the integrals in equation (63). It is assumed in equation (63) that the summation in ℓ extends only to L and that ϕ is of the form given by equation (64).

First consider the integrals over χ' from $-\pi$ to π in equation (63). The integrands can be put in the form of trigonometric polynomials of maximum degree $L + M - 1$. The following quadrature rule is exact if $f(\chi)$ is a trigonometric polynomial of degree less than or equal to $2M - 1$.

$$\int_{-\pi}^{\pi} f(\chi) d\chi = \frac{\pi}{M} \sum_{y=1}^{2M} f(\chi_y) \quad \text{where} \quad \chi_y = -\pi - \frac{\pi}{2M} + \frac{\pi}{M} y . \quad (66)$$

This rule is just the 2M-point composite midpoint rule over the

interval $[-\pi, \pi]$. Results similar to equation (66) can be found in reference (38).

Using this rule in equation (63), the integrations will be exact provided that $M \geq L$. In addition, the integrations over χ' are zero for $\ell > M - 1$ so L should be selected such that $L \geq M - 1$. Hence, it appears optimal to choose $L \leq M \leq L + 1$.

Consider the integrals over μ' in equation (63) from -1 to 1 . The integrands consist of products of two functions in μ' and at least one of the two functions is even. If one of the two functions in μ' is odd, the resulting integral is zero and can be dropped from the equation. Therefore, expressions containing odd functions in μ' can be dropped from equation (63). The remaining integrals over μ' have integrands that are even functions in μ' . In addition, the remaining integrations over μ' can be put in the following form:

$$\int_{-1}^1 (1 - (\mu')^2)^{n/2} d\mu' \quad \text{for } n = 0, 1, 2, \dots, N - 1 + L.$$

However, for properly chosen μ_s and w_s , $1 \leq s \leq N$,

$$\int_{-1}^1 (1 - (\mu')^2)^{n/2} d\mu' = \sum_{s=1}^N (1 - \mu_s^2)^{n/2} w_s \quad (67)$$

$$\text{for } n = 0, 1, 2, \dots, 2N - 1.$$

Thus, it appears that N should be chosen to be greater than or equal to L .

The above is not a standard numerical integration rule. Therefore, more details about selection of μ_s and w_s are presented. Consider

$$\int_{-1}^1 (1 - \mu^2)^{n/2} d\mu = \int_0^1 z^n g(z) dz = \sum_{s=1}^N z_s^n w_s$$

$$\text{where } z = (1 - \mu^2)^{\frac{1}{2}} \text{ and } g(z) = \frac{2z}{(1 - z^2)^{\frac{1}{2}}}.$$

Therefore, determination of μ_s and w_s for $1 \leq s \leq N$ corresponds to finding Gauss quadrature points and weights, z_s and w_s , for the interval $[0, 1]$ with nonnegative weight function $g(z)$.

Calculation of these points and weights can be done using well-developed procedures.³⁵ For example, if $N = 1$ then $w_1 = 2$ and to six decimal places, $\mu_1 = .618991$. If $N = 2$, then to six decimal places, $w_1 = .501092$, $w_2 = 1.49891$, $\mu_1 = .916788$, and $\mu_2 = .404704$.

Using these quadrature rules in equation (63), in which the summation in ℓ only extends to L , the following equation results.

$$\begin{aligned} K\phi(r, \mu, \chi) &= \sigma(r) \phi(r, \mu, \chi) - \sum_{\substack{\ell=0 \\ \ell \neq \text{odd}}}^L \frac{2\ell + 1}{4M} \sigma_\ell(r) \cdot \\ &\cdot \sum_{s=1}^N \sum_{y=1}^{2M} \phi(r, \mu_s, \chi_y) w_s P_\ell(\mu) P_\ell(\mu_s) \\ &- \sum_{\ell=0}^L \frac{2\ell + 1}{2M} \sigma_\ell(r) \sum_{s=1}^N \sum_{y=1}^{2M} \phi(r, \mu_s, \chi_y) w_s \cdot \end{aligned} \quad (68)$$

$$\sum_{\substack{v=1 \\ \ell+v \neq \text{odd}}}^{\ell} \frac{(\ell - v)!}{(\ell + v)!} P_{\ell}^v(\mu) P_{\ell}^v(\mu_s) \cos v(\chi - \chi_y) .$$

Note that in equation (68), $L \leq M \leq L + 1$ and $N \geq L$. A more precise relation between M and N was not found in the present investigation. Note, however, that the quadrature points, χ_y and μ_s for $1 \leq y \leq 2M$ and $1 \leq s \leq N$, completely determine the direction cosine point set. It appears that this set is different from those generally used in cylindrical geometry discrete ordinates methods although the equations derived in this section are also different from the equations used in other discrete ordinates methods for cylindrical geometry. Descriptions of point sets used in finite-difference discrete ordinates methods can be found in references (4) and (5). Notice also that $K\phi$ given by equation (68) is not equivalent to that given by equation (63) as the summation in ℓ was restricted to values of ℓ less than or equal to L . However, the terms in equation (63) for values of ℓ greater than L , assuming $L \leq M \leq L + 1$ and $N \geq L$, appear to be small in comparison with the terms for values of ℓ less than or equal to L which are exactly evaluated in equation (68). (In plane and spherical geometries, the terms for $\ell > L$ are identically zero in the analogous expressions for $K\phi$.)

Using the numerical integration rules just described, the inner products appearing in equation (62) can be exactly evaluated. The integrand is a polynomial of degree less than or equal to $2K - 2$ in r over each major subinterval. Therefore, a $K + 1$ -point Lobatto integration rule can be used for the integration over each major subinterval. The integrand in equation (62) is a trigonometric

polynomial of degree less than or equal to $2M - 1$ in χ and therefore the integration rule given by equation (66) can be used with $2M$ points. The integrand has functions of the form $(1 - \mu^2)^{n/2}$ where $n \leq 2N - 1$. Therefore, the N -point Gaussian quadrature method described by equation (67) can be used to evaluate the integrations with respect to μ .

Integrating by parts the first two terms in the inner product on the left hand side of equation (62) and using the above numerical integration rules, the following equation is obtained:

$$\sum_{t=1}^N w_t (1 - \mu_t^2)^{\frac{n+1}{2}} \sum_{p=0}^K \sum_{y=1}^{2M} w_p \frac{\pi}{M} L_{jim} \phi_i(r_{ip}, \mu_t, \chi_y) \quad (69)$$

$$= \sum_{t=1}^N w_t (1 - \mu_t^2)^{n/2} \sum_{p=0}^K \sum_{y=1}^{2M} w_p \frac{\pi}{M} r_{ip} \cos m\chi_y \cdot$$

$$\cdot \psi_{ji}(r_{ip}) [Q_i(r_{ip}, \mu_t, \chi_y) - K\phi_i(r_{ip}, \mu_t, \chi_y)]$$

where, in addition,

$$\phi_i(r_{ip}, \mu_t, \chi_y) = \phi_i(r_{ip}, \mu_t, -\chi_y) \quad (70)$$

$$\left\{ \begin{array}{ll} \phi_i(r_i, \mu_t, \chi_y) = \phi_{i-1}(r_i, \mu_t, \chi_y) & \text{for all } \chi_y \\ \phi_0(0, \mu_t, \chi_y) = \phi_0(0, \mu_t, \pi - |\chi_y|) & \text{for all } \chi_y \\ \phi_I(c, \mu_t, \chi_y) = 0 & \text{for } |\chi_y| > \pi/2 \end{array} \right. \quad (71)$$

and

$$\begin{aligned}
 L_{jim} \phi_i(r, \mu, \chi) = & -r \psi_{ji}'(r) \cos \chi \cos m\chi \phi_i(r, \mu, \chi) \\
 & + \frac{\cos \chi \cos m\chi}{r_{i+1} - r_i} [\phi_i(r_{i+1}, \mu, \chi) r_{i+1} \psi_{ji}(r_{i+1}) \\
 & - \phi_i(r_i, \mu, \chi) r_i \psi_{ji}(r_i)] - m \sin \chi \sin m\chi \psi_{ji}(r) \phi_i(r, \mu, \chi)
 \end{aligned}$$

for $0 \leq n \leq N - 1$, $1 \leq j \leq K - 1$, $0 \leq m \leq M - 1$,

and $0 \leq i \leq I - 1$.

Note that in equation (69), w_t for $1 \leq t \leq N$ refers to Gaussian quadrature weights and w_p for $0 \leq p \leq K$ refers to Lobatto quadrature weights. Equation (70) is obtained by symmetry considerations and the last two equations of (71) are obtained from the boundary conditions. The boundary at $r = 0$ is reflective and at $r = c$ a specified inward flux density is assumed. Note that the boundary conditions are only required to be satisfied at specified discrete directions.

Equations (69), (70), and (71) give $INM(K - 1)$ equations but including symmetry there are $INMK$ values of $\phi_i(r_{ip}, \mu_t, \chi_y)$ to be determined. However, for each pair of values of μ_t and χ_y , $\phi_i(r, \mu_t, \chi_y)$ is a polynomial of degree $K - 1$ in r over the major subinterval $[r_i, r_{i+1}]$. In equation (69), $\phi_i(r, \mu_t, \chi_y)$ is required at $K + 1$ points on this subinterval. Thus, $\phi_i(r, \mu_t, \chi_y)$

at one point, say r_{ic} , can be specified in terms of the other values of $\phi_i(r, \mu_t, \chi_y)$ in the major subinterval. Therefore, using Lagrange interpolation, equations identical in form to equations (60) and (61) are obtained. This procedure can be done for each pair of values of μ_t and χ_y for each major subinterval yielding the required additional INM equations.

As with plane or spherical geometry, it appears that the above equations can be solved in an iterative or noniterative manner. In addition, it appears likely that for values of N and M greater than two, the above equations will yield more accurate results than the diamond-difference discrete ordinates equations which contain an assumption of linearity between adjacent discrete directions. As for spherical geometry, a numerical investigation is required to compare the accuracy of this method with other discrete ordinates methods.

CHAPTER VII

NUMERICAL RESULTS

In the first section of this chapter, two numerical examples are presented in which the accuracy obtained using the standard multigroup equations is compared to that obtained using generalized multigroup equations based on two expansion functions in energy. Numerical results are obtained using the multigroup equations derived in Chapter IV with $M = 1$ and $M = 2$. These examples indicate that using two expansion functions in energy over each energy interval ($M = 2$) yields results as accurate as using the standard multigroup equations ($M = 1$) with eight times the number of energy intervals.

In the second section of this chapter, two computer codes based on the equations derived in Chapter V for plane geometry transport are described. One code is based on a linear approximation over spatial intervals ($K = 1$) and the other is based on a parabolic approximation over spatial intervals ($K = 2$). Both codes use a noniterative method of solution over each energy group. These computer codes are compared with each other and with a production version finite-difference discrete ordinates code. It is found that the noniterative method of solution used in these codes is faster than the iterative method used in the production-version discrete ordinates code. In addition, it is found in the example considered that the parabolic approximation is as accurate as the linear approximation while using only one-sixth the number of spatial intervals.

1. Comparisons of Multigroup Approximations

In this section, two physically reasonable examples are presented. Total flux densities are calculated for each example using the multigroup equations derived in Chapter IV. Either one or two expansion functions in energy are used over each energy interval, i.e. either $M = 1$ or $M = 2$. In each example, the exact solution is compared with results obtained using the standard multigroup equations ($M = 1$) and with results obtained using two expansion functions in energy over each energy interval ($M = 2$).

A. Example 1

In this example, a beam of neutrons is perpendicularly incident on an absorbing nonscattering infinite slab. The neutrons have energies ranging from 1 to 2 electron volts (eV).

For this example, the transport equation simplifies to:

$$\frac{\partial \psi}{\partial x}(x, E) + \sigma(E) \psi(x, E) = 0 \quad (72)$$

where $0 \leq x < \infty$, $1 \leq E \leq 2$

$\psi(x, E)$ = neutron flux density with neutron energy
E at position x in the slab

and $\sigma(E)$ = energy dependent total macroscopic cross
section (also absorption cross section in
this example) .

Assuming that $\psi(0, E) = \frac{1}{E \ln 2}$ and $\sigma(E) = E^{-\frac{1}{2}}$, the flux density at position x can be determined using equation (72). The exact total flux density at position x is then

$$\psi_T(x) = \int_1^2 \psi(x, E) dE = \frac{2}{\ln 2} \int_{\frac{x}{\sqrt{2}}}^x \frac{e^{-z}}{z} dz .$$

This problem is first solved numerically using the standard multigroup equations ($M = 1$) assuming that the neutron flux density approximation has the following form:

$$\psi^h(x, E) = b(E) \sum_{g=1}^G \phi_{1g}^h(x) \chi_g(E)$$

$$\text{where } b(E) = \frac{1}{E \ln 2}$$

$$\chi_g(E) = \begin{cases} 1 & \text{if } E_{g+1} \leq E \leq E_g \\ 0 & \text{otherwise} \end{cases}$$

$$\text{and } E_g = 2^{1/G} E_{g+1} \quad \text{where } E_{G+1} = 1 .$$

The energy group structure, i.e. energy partition, is based on equal logarithmic intervals. In other words, $\ln E_g - \ln E_{g+1} = \text{constant}$. This type of energy group structure is used extensively in practice and for a $1/E$ spectrum the total flux density is the same in each group. Note that the global approximation in the energy variable, $b(E)$, is assumed to have the same energy dependency as the incident

neutron flux density, $\psi(0, E)$. Also note that the approximate total flux density at position x is:

$$\psi_T^h(x) = \int_1^2 \psi^h(x, E) dE = \frac{1}{G} \sum_{g=1}^G \phi_{1g}^h(x) .$$

The method in Chapter IV is used to find a set of multigroup equations for equation (72). These equations are solved for various values of the total number of groups G . The following results are obtained for the approximate total neutron flux density as a function of position x for various values of G .

$$\text{For } G = 1 , \psi_T^h(x) = e^{-.84511x}$$

$$\text{For } G = 2 , \psi_T^h(x) = \frac{1}{2}(e^{-.91815x} + e^{-.77207x})$$

$$\text{For } G = 4 , \psi_T^h(x) = \frac{1}{4}(e^{-.95790x} + e^{-.87840x} + e^{-.80550x} + e^{-.73864x}) .$$

$$\begin{aligned} \text{For } G = 8 , \psi_T^h(x) = \frac{1}{8}(e^{-.97865x} + e^{-.93716x} + e^{-.89743x} + \\ + e^{-.85938x} + e^{-.82294x} + e^{-.78805x} + \\ + e^{-.75464x} + e^{-.72265x}) . \end{aligned}$$

Total flux densities calculated using the above equations for various positions into the slab are given in Table 1. Note that since $\frac{1}{\sqrt{2}} \leq \sigma \leq 1$, the results are given for up to 22.6 to 32 mean free paths into the slab which corresponds to quite a deep penetration.

TABLE 1
COMPARISON OF RESULTS FOR EXAMPLE 1

Position x	Total Flux Density					
	Exact	M = 1, G = 1	M = 1, G = 2	M = 1, G = 4	M = 1, G = 8	M = 2, G = 1
1	.43104	.42951	.43066	.43094	.43101	.43102
2	.18710	.18448	.18645	.18694	.18706	.18709
4	.03598	.03403	.03550	.03586	.03595	.03598
8	.001434	.001158	.001362	.001415	.001429	.001435
16	2.86×10^{-6}	1.34×10^{-6}	2.37×10^{-6}	2.73×10^{-6}	2.82×10^{-6}	2.82×10^{-6}
24	6.86×10^{-9}	1.55×10^{-9}	4.62×10^{-9}	6.21×10^{-9}	6.69×10^{-9}	6.29×10^{-9}
32	1.82×10^{-11}	1.80×10^{-12}	9.40×10^{-12}	1.53×10^{-11}	1.74×10^{-11}	1.46×10^{-11}

This problem is now solved numerically using two expansion functions in energy over each energy interval ($M = 2$) but for only one energy group or $G = 1$. In this problem it is reasonable to let $b_{11}(E) = 1$ and $b_{21}(E) = E$. The flux density approximation then has the form

$$\psi^h(x, E) = b(E) (\phi_{11}^h(x) + E \phi_{21}^h(x)) \quad \text{for } 1 \leq E \leq 2$$

$$\text{where } b(E) = \frac{1}{E \ln 2}.$$

The total flux density is then

$$\psi_T^h(x) = \phi_{11}^h(x) + \phi_{21}^h(x)/\ln 2.$$

The method in Chapter IV is again used to find a set of multi-group equations for equation (72) assuming $M = 2$ and $G = 1$. Using this approach, the total flux density approximation as a function of position x is found to be:

$$\psi_T^h(x) = .47007 e^{-.756228x} + .52993 e^{-.923955x}.$$

This function is tabulated in Table 1 for various values of x .

As can be seen in Table 1, the flux density values for $M = 2$, $G = 1$ are more accurate than those for $M = 1$, $G = 8$ up to a distance of $x = 16$ (11.3 to 16 mean free paths). For greater distances, the standard multigroup approximation using eight energy

intervals ($M = 1, G = 8$) is more accurate than using two expansion functions in energy over one energy interval ($M = 2, G = 1$).

B. Example 2

In this example, a fission source is present uniformly in an infinite medium of homogeneous material. The quantity of interest is the total flux density from .2 to $\sqrt{2}$ million electron volts (MeV).

For an infinite homogeneous medium, the neutron transport equation simplifies to

$$\sigma(E) \psi(E) = \int_B \psi(E') \sigma f(E' \rightarrow E) dE' + Q(E).$$

$$\text{It is assumed that } \sigma f(E' \rightarrow E) = \begin{cases} \frac{1}{E'} & \text{for } 0 \leq E \leq E' \\ 0 & \text{otherwise.} \end{cases} \quad \text{This is a}$$

physically reasonable energy-transfer cross section for homogeneous

material. Note that $\sigma_s(E') = \int_B \sigma f(E' \rightarrow E) dE = 1$ where $\sigma_s(E')$ is

the scattering cross section at energy E' . (Actually, from .2 MeV to 10 MeV, the hydrogen scattering cross section decreases about a factor of 10 but to simplify this example it is assumed to be constant.)

It is now assumed that $\sigma(E) = \sigma_a(E) + \sigma_s(E) = \sigma_a + 1 = \sigma$ where σ_a is the absorption cross section and σ is the total cross section. Two cases will be studied in this example, In the first

case, σ_a is zero which corresponds with no absorption and therefore $\sigma = 1$. In the second case, σ_a is one and therefore $\sigma = 2$.

The fission neutron spectrum, $Q(E)$ has the following approximate form:³⁹

$$Q(E) \approx .242 e^{-E} (e^{\sqrt{2E}} - e^{-\sqrt{2E}}) .$$

Very few neutrons originate with energies greater than 10 MeV (about .1%) so the total energy interval considered is [.2 MeV, 10 MeV] .

Using the above assumptions, the transport equation for this example is:

$$\sigma\psi(E) = \int_E^{10} \psi(E') \frac{1}{E'} dE' + Q(E) \quad (73)$$

for $.2 \leq E \leq 10$

where $\sigma = 1$ or $\sigma = 2$.

The exact solution to equation (73) is:

$$\psi(E) = \frac{Q(E)}{\sigma} + E^\beta \int_E^{10} \frac{Q(x) x^\alpha}{\sigma^2} dx$$

where $\beta = -\frac{1}{\sigma}$ and $\alpha = \frac{1-\sigma}{\sigma}$.

The quantity of interest in this example is the total flux density in the energy range $.2 \leq E \leq \sqrt{2}$ MeV . The exact total flux density

in this energy range is:

$$\begin{aligned} \psi_T = \int_{.2}^{\sqrt{2}} \psi(E) dE &= \frac{1}{\sigma} \int_{.2}^{\sqrt{2}} Q(E) dE + \int_{.2}^{\sqrt{2}} \frac{Q(E) E^\alpha}{\sigma^2} \int_{.2}^E x^\beta dx dE \\ &+ \int_{.2}^{\sqrt{2}} E^\beta dE \int_{\sqrt{2}}^{10} \frac{Q(E) E^\alpha}{\sigma^2} dE . \end{aligned}$$

When $\sigma = 1$, the exact total flux density has the value 2.0148.

When $\sigma = 2$, the total flux density has the value .42359.

This problem is first solved numerically for $M = 1$ (standard multigroup calculation) using from 2 to 16 energy intervals. For this case, the function, $b(E)$, is assigned the fission neutron spectrum. The approximation to the flux density then has the following form:

$$\psi^h(E) = b(E) \sum_{g=1}^G \phi_{1g}^h \chi_g(E)$$

where $b(E) = Q(E)$, the fission neutron spectrum. As in the first example, equal logarithmic energy intervals are used. The approximate total flux density from .2 to $\sqrt{2}$ MeV is:

$$\psi_T^h = \int_{.2}^{\sqrt{2}} \psi^h(E) dE = \sum_{g=\frac{G}{2}+1}^G c_g \phi_{1g}^h$$

$$\text{where } c_g = \int_{E_{g+1}}^{E_g} b(E) dE \text{ and } G \text{ is an even integer.}$$

Using the multigroup equations derived in Chapter IV, a system of equations is obtained for calculating ϕ_{1g}^h , $1 \leq g \leq G$. Values

of ψ_T^h were then calculated as the number of energy intervals, G , ranged from 2 to 16. The results obtained for $\sigma = 1$ and $\sigma = 2$ are given in Table 2.

TABLE 2
COMPARISON OF RESULTS FOR EXAMPLE 2

σ	Total Flux Density (from .2 to $\sqrt{2}$ MeV)					
	Exact	M=1,G=2	M=1,G=4	M=1,G=8	M=1,G=16	M=2,G=2
1	2.0148	2.7934	2.1935	2.0586	2.0257	1.9970
2	.42359	.43942	.42763	.42461	.42384	.42358

Next the problem is solved numerically for $M = 2$ (two expansion functions in energy over each energy interval) and two energy groups. For this case, $b(E)$ is again assigned the fission neutron spectrum, $b_{1g}(E)$ is assigned the value unity, and $b_{2g}(E)$ is assigned the function $1/E$. This function was selected for $b_{2g}(E)$ since the flux density tends to a $1/E$ distribution in energy in the presence of a scatterer such as hydrogen.²³ The approximation to the flux density for this case then has the form:

$$\psi^h(E) = b(E) \sum_{g=1}^2 x_g(E) \left(\phi_{1g}^h + \frac{1}{E} \phi_{2g}^h \right)$$

where $b(E) = Q(E)$.

For this case, $E_1 = 10$ MeV, $E_2 = \sqrt{2}$ MeV, and $E_3 = .2$ MeV.

The total flux density from .2 to $\sqrt{2}$ MeV in this case is:

$$\psi_T^h = \phi_{12}^h \int_{.2}^{\sqrt{2}} b(E) dE + \phi_{22}^h \int_{.2}^{\sqrt{2}} \frac{b(E)}{E} dE .$$

Using the equations derived in Chapter IV, a system of equations is obtained for ϕ_{1g}^h and ϕ_{2g}^h for $g = 1, 2$. After solving this system, ψ_T^h is calculated using the above equation. The results are given in Table 2 for values of σ of 1 and 2. As can be seen in Table 2, the results for $M = 2, G = 2$ are approximately as accurate as those for $M = 1, G = 16$.

Summarizing both examples, a two term expansion in energy over each energy interval ($M = 2$) generally gave results as accurate as a one term expansion (standard multigroup calculation) using eight times the number of energy intervals.

2. Comparisons of Approximations for Plane Geometry

In this section, two computer codes are described which are based on the equations derived in Chapter V, specifically equations (44) and (45), for plane geometry transport. In each computer code, the value of M used is one so the codes are based on the standard multigroup equations. However, in the first code, Code A, a linear approximation over spatial intervals is used ($K = 1$) and in the second code, Code B, a parabolic approximation over spatial intervals is used ($K = 2$). Both codes use a noniterative method of solution over each energy group. Codes A and B are designed to perform criticality calculations, i.e. to find the multiplication factor or eigenvalue,

for bare infinite slabs. Computational results using these codes are compared with those obtained using ANISN-ORNL⁴⁰ which is a production-version finite-difference discrete ordinates code.

It was shown in Chapter V that for $K = 1$ and $M = 1$, equations (44) and (45) reduce to the diamond-difference discrete ordinates equations. Therefore, Code A is based on the diamond-difference discrete ordinates equations. The difference between Code A and production-version diamond-difference discrete ordinates computer codes is that Code A uses a noniterative method of solution for the flux densities in each energy group whereas production discrete ordinates codes use an iterative method of solution. Code B also differs from existing discrete ordinates computer codes. Code B uses a parabolic approximation over spatial intervals as well as a noniterative method of solution. The method nearest to that used in Code B appears to be the fourth-order finite-difference method which Neta and Victory¹³ describe as being a quadratic continuous method. Another similar method is that used in ONETRAN¹⁴ which Hennart¹¹ describes as being a continuous piecewise parabolic approximation. Code B differs from other codes that are based on parabolic approximations by using a non-iterative method of solution over each energy interval.

In conventional neutron transport discrete ordinates computer codes, the flux density in each energy group is solved in an iterative manner over positions and directions. These iterations are called inner iterations. After a specified number of inner iterations or after some convergence criteria are satisfied, the inner iterations for that group are terminated and the next group, of lower energy, is

then iterated on a similar manner. After all groups have been iterated on in this fashion, if upscattering or fission is present, the entire procedure is repeated over all groups. Each iteration over all groups is called an outer iteration. Outer iterations cease when specified convergence criteria are met.

Codes A and B differ from conventional computer codes in that there are no inner iterations. A matrix is formulated for each energy group which has the form of matrix A described in Chapter V. For each energy group, the matrix is decomposed into the product of a lower triangular and an upper triangular matrix. The flux densities for each direction and position are then rapidly calculated by the process of forward and backward substitution. Since the matrices are banded, storage requirements are kept low. Although a direct method replaces inner iterations, outer iterations are still used.

The production-version one-dimensional geometry discrete ordinates code, ANISN-ORNL, was obtained for comparison purposes. ANISN-ORNL was obtained from the Radiation Shielding Information Center located at Oak Ridge, Tennessee. Subsequently, the code was adapted to The University of Tennessee's IBM 360 computer system by University of Tennessee computer consultant Sue Smith.

For use in calculations to compare Codes A and B with ANISN-ORNL, a set of six-group cross sections was obtained by spectrum weighting a set of 218-group cross sections that are described in reference (41). The seven energy boundaries of the six groups in units of electron volts are: 1.00×10^{-5} , 1.25×10^{-1} , 3.05, 5.5×10^2 , 1.00×10^5 , 1.85×10^6 , and 2.00×10^7 from lowest

energy to highest energy, respectively. Groups 5 and 6 are thermal groups and a nonzero upscattering cross section from group 6 to group 5 is present. In the cross section set, the degree of the scattering expansion was selected to be two. In other words, L is equal to two in the cross section set.

For comparison purposes, Code A, Code B, and ANISN-ORNL were each used to calculate the multiplication factor or eigenvalue for a bare infinite slab of a water-uranium mixture containing .05g/cc uranium-235 and .00365g/cc uranium-238. The thickness of the slab was selected to be 18cm after preliminary calculations indicated that this thickness would yield a calculated multiplication factor close to unity using the six-group cross section set. This thickness is within 20% of that indicated in reference (42) for a critical bare slab of the above composition. All calculations modeled the bare slab using a reflecting boundary condition to represent the center of the slab and a nonreentrant boundary condition to represent one side. The thickness of the modeled slab was therefore one-half of the real slab or 9cm. In each calculation, the number of directions, N , was four and the number of groups, G , was six.

In these calculations, Code A was first compared with ANISN-ORNL. ANISN-ORNL was run using the diamond-difference equations with a maximum of 10 inner iterations for each group in each outer iteration. In each outer iteration for Code A, the flux densities in groups 5 and 6 were calculated a second time to improve their accuracy since group 6 upscattered to group 5. For both codes the slab was divided into 8 equal intervals and the calculations were continued for 10

outer iterations. The calculated eigenvalues for even outer iterations are given in Table 3. Eventually, the same eigenvalue is obtained since both codes are based on the same discrete ordinates equations. After 10 outer iterations, the maximum flux deviation between iterations 9 and 10 was 8.03×10^{-6} for the ANISN-ORNL calculation compared with 7.75×10^{-7} for the Code A calculation. In addition, the time required by ANISN-ORNL was 6.02 seconds compared with 5.34 seconds for Code A. Therefore, Code A was faster and more accurate than ANISN-ORNL. Although the ANISN-ORNL calculation was possibly not optimized, Code A, being an experimental program was not optimized either. In addition, the same compiler was used to compile both programs to eliminate time differences due to different compilers.

TABLE 3
COMPARISON BETWEEN ANISN-ORNL AND CODE A

Outer Iteration	Calculated Eigenvalue (Multiplication Factor)	
	ANISN-ORNL	Code A
2	.97144	.98369
4	1.0026	1.0047
6	1.0051	1.0052
8	1.0052	1.0052
10	1.0052	1.0052

Next, Code B was compared with Code A. Using Code A, the eigenvalue of the above uranium-water mixture was calculated for various

numbers of intervals. The eigenvalue of this system was also calculated using Code B with four intervals. The results of these eigenvalue calculations are given in Table 4. The best estimate of the actual eigenvalue of the system using Code A is 1.0066 as the accuracy of the calculated eigenvalue is expected to improve as the number of intervals increases. Therefore, Code B using only 4 intervals was at least as accurate as Code A using 24 intervals. In addition, Code B required 3.24 seconds to do 10 outer iterations and the maximum flux deviation between the last two outer iterations was 8.94×10^{-7} .

TABLE 4
COMPARISON BETWEEN CODE A AND CODE B

Number of Intervals	Calculated Eigenvalue	
	Code A	Code B
4	1.0007	1.0067
8	1.0052	
16	1.0062	
24	1.0065	
32	1.0066	

In summary, the results presented in this section indicate that inner iterations can be effectively replaced by a noniterative procedure. The storage requirements may be somewhat greater than for

discrete ordinates iterative methods but they are kept low since the matrices are banded. One disadvantage of a noniterative procedure may be the difficulty in performing a flux fix-up if a calculated flux density value is negative. However, for a parabolic approximation, as that used in Code B, the increased accuracy may substantially reduce the possibility of obtaining negative flux densities. In addition, a noniterative method may be computationally advantageous using vector processing computers.

An interesting undertaking would be to write and test a computer code that uses perhaps $M = 2$ and $K = 3$. Based on the results of the last two sections, it appears that such a code would be fast and accurate.

CHAPTER VIII

SUMMARY AND CONCLUSIONS

In the present investigation, Galerkin's method was applied to the steady-state energy-dependent neutron transport equation. In this method, the angular flux density, $\psi(\vec{r}, \vec{\Omega}, E)$, was defined to be the product $b(E) \phi(\vec{r}, \vec{\Omega}, E)$ where $b(E)$ is a specified global approximation with respect to energy for the angular flux density. For example, the function $b(E)$ may be chosen to be the Maxwellian- $1/E$ -fission energy spectrum. Using Galerkin's method, $\phi(\vec{r}, \vec{\Omega}, E)$ was then approximated. This approximation consisted of an expansion of linearly independent functions in the variables $\vec{r}$, $\vec{\Omega}$, and E .

The energy variable was treated separately from the direction and position variables using Galerkin's method. This approach yielded a generalized set of multigroup neutron transport equations. These multigroup equations reduced to the standard multigroup equations for a single expansion function in energy over each energy interval. (In fact, this was the purpose for the presence of the global approximation with respect to energy, $b(E)$.) However, numerical examples indicated that using two expansion functions with respect to energy over each energy interval gives results as accurate as using the standard multigroup equations but with eight times the number of energy intervals. Disadvantages of using more than one expansion function in energy over each energy interval are that additional spectrum-weighted cross-sections are required and additional computations for each energy group in the transport calculations are required.

In addition, an error analysis was performed in which the energy variable was treated separately from the position and direction variables. It was estimated that the error in the approximate solution is proportional to the maximum energy interval width for the standard multigroup equations and proportional to the square of the maximum energy interval width for a linear expansion over energy intervals.

Next, using a Galerkin method adapted from reference (34) for solving first-order ordinary differential equations, equations relating flux densities at specified positions and directions were derived for one-dimensional plane, spherical, and cylindrical geometries. The aim of this part of the present investigation was to derive equations that could be solved iteratively or noniteratively for multigroup flux densities with an accuracy better than that of well-established numerical methods for solving the multigroup transport equations. In deriving these equations, the expansion functions and numerical integration rules were selected so that the inner products resulting in Galerkin's method would be evaluated exactly. In addition, the flux density approximation was assumed to be a continuous piecewise polynomial with respect to the position variable.

In plane geometry, the Galerkin method with a linear approximation in the position variable led to equations identical to the finite-difference discrete ordinates equations. In addition, the Galerkin method using a parabolic approximation in the position variable appears to give the same equations as those obtained in the fourth-order finite-difference method described in reference

(13). It was demonstrated that these equations can be efficiently, i.e. with low storage requirements, solved noniterative over each energy group with computer times less than those required by ANISN-ORNL,⁴⁰ a production-version finite-difference code. In addition, in a numerical comparison, the parabolic approximation using 4 spatial intervals gave results as accurate as the linear approximation using 24 spatial intervals.

In spherical and cylindrical geometries, the Galerkin method did not yield equations that were equivalent to the finite-difference discrete ordinates equations for any order of approximation. Although numerical comparisons were not made, it is reasonable to suppose that the equations derived would yield more accurate results than the finite-difference discrete ordinates equations. In the derivation of the finite-difference (diamond-difference) discrete ordinates equations for spherical and cylindrical geometries, the flux density is assumed to be linear between adjacent discrete directions. However, in the equations derived in the present investigation, the polynomial approximation with respect to the direction variables is not restricted to being linear.

LIST OF REFERENCES

LIST OF REFERENCES

1. R.A. Ackroyd, "The Why and How of Finite Elements", Annals of Nuclear Energy, 8, pp. 539-566 (1981).
2. R. Sanchez and N.J. McCormick, "A Review of Neutron Transport Approximations," Nuclear Science and Engineering, 80, pp.481-535 (1982).
3. N.M. Schaeffer, editor, Reactor Shielding for Nuclear Engineers, Published by U.S. Atomic Energy Commission, TID-25951 (1973).
4. C.E. Lee, "The Discrete S_n Approximation to Transport Theory," LA-2595, Los Alamos Scientific Laboratory (1962).
5. B.G. Carlson and K.D. Lathrop, "Transport Theory - The Method of Discrete Ordinates," Computing Methods in Reactor Physics, edited by H. Greenspan, C.N. Kelber, and D. Okrent, Gordon and Breach, New York (1968).
6. E.M. Gelbard, J.A. Davis, L.A. Hageman, "Solution of the Discrete Ordinate Equations in One and Two Dimensions," Proceedings of a Symposium in Applied Mathematics of the American Mathematical Society and The Society for Industrial and Applied Mathematics, Vol. I, SIAM-AMS Proceedings, Transport Theory, American Mathematical Society, Providence, Rhode Island (1969).
7. S.M. Lee and R. Vaidyanathan, "Comparison of the Order of Approximation in Several Spatial Difference Schemes for the Discrete-Ordinates Transport Equation in One-Dimensional Plane Geometry," Nuclear Science and Engineering, 76, pp. 1-9 (1980).
8. E.W. Larsen and W.F. Miller, Jr., "Convergence Rates of Spatial Difference Equations for the Discrete Ordinates Neutron Transport Equations in Slab Geometry," Nuclear Science and Engineering, 73, pp. 76-83 (1980).
9. E.W. Larsen and P. Nelson, "Finite-Difference Approximations and Superconvergence for the Discrete-Ordinate Equations in Slab Geometry," SIAM Journal of Numerical Analysis, 19, pp. 334-348 (1982).
10. R.E. Alcouff, E.W. Larsen, W.F. Miller, Jr., and B.R. Wienke, "Computational Efficiency of Numerical Methods for the Multigroup, Discrete-Ordinates Neutron Transport Equations: The Slab Geometry Case," Nuclear Science and Engineering, 71, pp. 111-127 (1979).
11. J.P. Hennart, "A Unified Formalism for Spatial Discretization Schemes of Transport Equations in Slab Geometry," Annals of Nuclear Energy, 8, pp. 677-681 (1981).

12. P. Barbucci and F. Di Pasquantonio, "Computational Efficiency of Some Numerical Methods for the Multigroup Discrete Ordinates Solution of Stationary Boltzmann Equations for Neutrons and Gamma Rays in Slab Geometry," Nuclear Science and Engineering, 82, pp. 448-475 (1982).
13. B. Neta and H.D. Victory, Jr., "A New Fourth-Order Finite-Difference Method for Solving Discrete-Ordinates Slab Transport Equations," SIAM Journal of Numerical Analysis, 20, pp. 94-105 (1983).
14. T.R. Hill, "ONETRAN: A Discrete Ordinates Finite Element Code for the Solution of the One-Dimensional Multigroup Transport Equation," LA-5990-MS, Los Alamos Scientific Laboratory (1975).
15. W.R. Martin and J.J. Duderstadt, "Finite Element Solutions of the Neutron Transport Equation with Applications to Strong Heterogeneities," Nuclear Science and Engineering 62, pp. 371-390 (1977).
16. W.R. Martin, C.E. Yehner, L. Lorence, and J.J. Duderstadt, "Phase-Space Finite Element Methods Applied to the First-Order Form of the Transport Equation," Annals of Nuclear Energy 8, pp. 633-646 (1981).
17. R.T. Ackroyd, "A Finite Element Method for Neutron Transport - I. Some Theoretical Considerations," Annals of Nuclear Energy, 5, pp. 75-94 (1978).
18. J. Galliara and M.M.R. Williams, "A Finite Element Method for Neutron Transport - II. Some Practical Considerations," Annals of Nuclear Energy, 6, pp. 205-223 (1979).
19. R.T. Ackroyd, A.K. Ziver, A.J.H. Goddard, "A Finite Element Method for Neutron Transport. Part IV: A Comparison of Some Finite Element Solutions of Two Group Benchmark Problems with Conventional Solutions," Annals of Nuclear Engineering, 7, pp. 335-349 (1980).
20. J. Gerin-Roze and P. Lesaint, "Isoparametric Finite Element Methods for Two-Dimensional Transport Calculations," International Journal for Numerical Methods in Engineering, 10, pp. 171-183 (1976).
21. P. Lesaint and P.A. Raviart, "On a Finite Element Method for Solving the Neutron Transport Equation," Mathematical Aspects of Finite Elements in Partial Differential Equations, edited by Carl de Boor, Academic Press, New York (1974).
22. S. Ukai, "Solution of Multi-Dimensional Neutron Transport Equation by Finite Element Method," Journal of Nuclear Science and Technology, 9, pp. 366-373 (1972).

23. G.I. Bell and S. Glasstone, Nuclear Reactor Theory, Van Nostrand Reinhold Company, New York (1970).
24. L.N. Slobodeckii, "Generalized Sobolev Spaces and Their Application to Boundary Problems for Partial Differential Equations," Transactions American Mathematical Society, Series 2, 57, pp. 207-273 (1966).
25. P.G. Ciarlet, The Finite Element Method for Elliptic Problems, North Holland Publishing Company, New York (1980).
26. S. Agmon, Lectures on Elliptic Boundary Value Problems, D. Van Nostrand Company, Inc., Princeton, New Jersey (1965).
27. A.R. Mitchell and R. Wait, The Finite Element Method in Partial Differential Equations, John Wiley and Sons, London (1977).
28. G. Strang and G.J. Fix, An Analysis of the Finite Element Method, Prentice-Hall, Inc., Englewood Cliffs, New Jersey (1973).
29. V.A. Dougalis and M.D. Gunzburger, Lecture Notes on the Finite Element Method, Unpublished, University of Tennessee, Knoxville, Tennessee (1979).
30. J.J. Duderstadt and W.R. Martin, Transport Theory, John Wiley and Sons, New York (1979).
31. W.W. Engle, Jr. and F.R. Mynatt, "A Comparison of Two Methods of Inner Iteration Convergence Acceleration in Discrete Ordinates Codes," Transactions of the American Nuclear Society, 2, No. 1, pp. 193-194 (1968).
32. R.E. Alcouffe, "Diffusion Synthetic Acceleration Methods for the Diamond-Differenced Discrete-Ordinates Equations," Nuclear Science and Engineering, 64, pp. 344-355 (1977).
33. A. Jennings, Matrix Computation for Engineers and Scientists, John Wiley and Sons, New York (1977).
34. B.L. Hulme, "Discrete Galerkin and Related One-Step Methods for Ordinary Differential Equations," Mathematics of Computation, 26, Number 120, pp. 881-891 (1972).
35. P.J. Davis and P. Rabinowitz, Methods of Numerical Integration, Academic Press, New York (1975).
36. F.B. Hildebrand, Introduction to Numerical Analysis, 2nd Edition, McGraw-Hill, Inc., New York (1974).
37. B. Wendroff, Theoretical Numerical Analysis, Academic Press, Inc., New York (1966).

38. D.M. Young, R.T. Gregory, A Survey of Numerical Mathematics, Vol. 1, Addison-Wesley Publishing Company, Reading, Massachusetts (1972).
39. H. Etherington, editor, Nuclear Engineering Handbook, McGraw-Hill, Inc., New York (1958).
40. W.W. Engle, Jr., "ANISN-ORNL: A One-Dimensional Discrete Ordinates Transport Code with Anisotropic Scattering," RSIC Code Package CCC-254 (1968, updated 1975).
41. W.E. Ford, III, C.C. Webster, and R.M. Westfall, "A 218-Group Neutron Cross-section Library in the AMPX Master Interface Format for Criticality Studies," ORNL/CSD/TM-4 (1976).
42. Reactor Physics Constants, prepared by Argonne National Laboratory, ANL 5800, 2nd edition, page 186 (1963).

VITA

Edward James Allen was born in Columbus, Wisconsin on July 23, 1949. He graduated from Portage High School in Portage, Wisconsin in 1967 and entered the University of Wisconsin at Madison the same year. In August 1971, Edward graduated from the University of Wisconsin with a Bachelor of Science degree in Nuclear Engineering. He obtained a Master of Science degree in Nuclear Engineering at the University of Wisconsin in December, 1972.

Edward then worked at the University of Wisconsin in the Nuclear Engineering Department for one year. In January 1974, he was hired by Oak Ridge National Laboratory where he worked in the Engineering Technology Division on a variety of nuclear engineering projects until July, 1980. During this time, he obtained his Professional Engineering certification.

Edward resumed graduate studies at The University of Tennessee in September, 1979, working toward a Doctor of Philosophy degree in mathematics specializing in numerical analysis. The Doctor of Philosophy degree was awarded to him in August, 1983.

Edward married the former Linda Joy Svoboda on September 20, 1980. Linda and Edward are currently employed in the Mathematics Department at the University of North Carolina at Asheville.