

*Formulación y Aplicación del Método  
de Elementos Finitos Mínimos  
Cuadrados a un Problema de  
Dispersión de Onda y Comparación  
con otros Métodos Numéricos*

Carlos E. Cadenas R.

3 de marzo de 2003

UNIVERSIDAD DE CARABOBO  
FACULTAD EXPERIMENTAL DE CIENCIAS Y TECNOLOGÍA  
DEPARTAMENTO DE MATEMÁTICAS

*Formulación y Aplicación del Método de Elementos Finitos  
Mínimos Cuadrados a un Problema de Dispersión de Onda y  
Comparación con otros Métodos Numéricos*

Trabajo presentado por el profesor Carlos E. Cadenas R. ante la ilustre  
Universidad de Carabobo como credencial de mérito para ascender en el  
escalafón a la Categoría de Profesor Agregado.

Bárbula, 3 de marzo de 2003

## Resumen

En esta investigación se desarrollan formulaciones de varios métodos numéricos para hacer un estudio comparativo de los mismos basándose en la aplicación a un problema de dispersión de ondas en una dimensión, el cual viene dado por la ecuación de Helmholtz, con la finalidad de conocer a profundidad las ventajas y desventajas del método de Elementos Finitos Mínimos Cuadrados. Los métodos a comparar con éste son Diferencias Finitas y Elementos Finitos Galerkin Mixto. La comparación se basa fundamentalmente en el estudio del orden de convergencia de dichos métodos y el análisis de la contaminación basada en la dispersión de la solución numérica. También se estudia el efecto del incremento de la frecuencia  $k$  en todos los métodos presentados. Se presentan varios ejemplos que validan las conclusiones.

# Dedicatoria

*A mis padres.*

## **Agradecimiento**

Al profesor Vianey Villamizar por su guía y colaboración en la realización de este trabajo.

**Muchísimas Gracias!**

Carlos E. Cadenas R.

# Índice general

|                                                                |          |
|----------------------------------------------------------------|----------|
| <b>1. PLANTEAMIENTO DEL PROBLEMA</b>                           | <b>1</b> |
| 1.1. Formulación del Problema de Dispersión de Ondas en 1-D    | 1        |
| 1.2. Objetivo General . . . . .                                | 2        |
| 1.3. Objetivos Específicos . . . . .                           | 3        |
| 1.4. Enfoque Metodológico . . . . .                            | 3        |
| 1.5. Área de Conocimiento . . . . .                            | 4        |
| 1.6. Solución Analítica . . . . .                              | 5        |
| 1.7. Sistema de Ecuaciones Diferenciales de Primer Orden . .   | 5        |
| <b>2. FORMULACIÓN GENERAL DE LOS MÉTODOS NUMÉRICOS</b>         | <b>7</b> |
| 2.1. Introducción del Método de Diferencias Finitas . . . . .  | 8        |
| 2.2. Formulación del Método de Diferencias Finitas Implícito . | 10       |

|                                                                                                              |           |
|--------------------------------------------------------------------------------------------------------------|-----------|
| 2.3. Introducción del Método de Elementos Finitos . . . . .                                                  | 10        |
| 2.3.1. Consideraciones Comunes . . . . .                                                                     | 11        |
| 2.4. Método de Elementos Finitos . . . . .                                                                   | 14        |
| 2.4.1. Construcción de funciones de Forma . . . . .                                                          | 15        |
| 2.4.2. Formulación del Método de Elementos Finitos Galerkin Mixto . . . . .                                  | 21        |
| 2.4.3. Formulación del Método de Elementos Finitos Mínimos Cuadrados para Sistemas de Primer Orden . . . . . | 23        |
| <b>3. APLICACIÓN DE LOS MÉTODOS NUMÉRICOS A LA ECUACIÓN DE HELMHOLTZ EN 1-D</b>                              | <b>25</b> |
| 3.1. Método de Diferencias Finitas Implícito . . . . .                                                       | 25        |
| 3.2. Método de Elementos Finitos Galerkin Mixto . . . . .                                                    | 27        |
| 3.3. Método de Elementos Finitos Mínimos Cuadrados para Sistemas de Primer Orden . . . . .                   | 30        |
| <b>4. Resultados y Conclusiones</b>                                                                          | <b>40</b> |
| 4.1. Análisis del Orden de Convergencia . . . . .                                                            | 40        |
| 4.1.1. Método de Diferencias Finitas Implícito . . . . .                                                     | 41        |
| 4.1.2. Método de Elementos Finitos Galerkin Mixto . . . . .                                                  | 41        |

|                                                                                                      |    |
|------------------------------------------------------------------------------------------------------|----|
| 4.1.3. Método de Elementos Finitos Mínimos Cuadrados<br>para Sistemas de Primer Orden . . . . .      | 50 |
| 4.2. Análisis de la Dispersión . . . . .                                                             | 55 |
| 4.2.1. Método de Diferencias Finitas Implícito . . . . .                                             | 59 |
| 4.2.2. Método de Elementos Finitos Galerkin Mixto . . . .                                            | 62 |
| 4.2.3. Método de Elementos Finitos Mínimos Cuadrados<br>para Sistemas de Primer Orden . . . . .      | 66 |
| 4.3. Comentarios Finales y Conclusiones . . . . .                                                    | 68 |
| 4.4. Trabajos futuros . . . . .                                                                      | 69 |
| A. Programación del Método de Diferencias Finitas Implícito                                          | 71 |
| B. Programación del Método de Elementos Finitos Galerkin Mixto                                       | 75 |
| C. Programación del Método de Elementos Finitos Mínimos Cua-<br>drados para Sistemas de Primer Orden | 80 |
| D. Cálculo de la Dispersión para el Método de Diferencias Finitas<br>Implícito                       | 85 |
| D.1. Ecuación en Diferencias . . . . .                                                               | 85 |
| D.2. Programa en Maple Para Resolver la Ecuación en Diferen-<br>cias . . . . .                       | 86 |

|                                                                                                                 |     |
|-----------------------------------------------------------------------------------------------------------------|-----|
| E. Cálculo de la Dispersión para el Método de Elementos Finitos Galerkin Mixto                                  | 87  |
| E.1. Programa en Maple Para Obtener la Ecuación en Diferencias                                                  | 87  |
| E.2. Programa en Maple Para Resolver la Ecuación en Diferencias . . . . .                                       | 88  |
| F. Cálculo de la Dispersión para el Método de Elementos Finitos Mínimos Cuadrados para Sistemas de Primer Orden | 90  |
| F.1. Programa en Maple Para Obtener la Ecuación en Diferencias                                                  | 90  |
| F.2. Programa en Maple Para Resolver la Ecuación en Diferencias. Versión I . . . . .                            | 91  |
| F.3. Programa en Maple Para Resolver la Ecuación en Diferencias. Versión II . . . . .                           | 93  |
| G. Cálculo de las Integrales Involucradas en los Métodos                                                        | 94  |
| H. Método de Diferencias Finitas Dominio Tiempo Explícito                                                       | 96  |
| H.1. Formulación del Método de Diferencias Finitas Dominio Tiempo Explícito . . . . .                           | 96  |
| H.2. Método de Diferencias Finitas Dominio Tiempo Explícito                                                     | 97  |
| I. Programación del Método de Diferencias Finitas Dominio Tiempo Explícito                                      | 101 |

|                                                                                                                              |            |
|------------------------------------------------------------------------------------------------------------------------------|------------|
| <b>J. Formulación y Aplicación del Método de Elemento Finito Mínimos Cuadrados a un Problema de Dispersión de Onda en 2D</b> | <b>104</b> |
| J.1. Ecuación de Helmholtz en 2D . . . . .                                                                                   | 105        |
| J.1.1. Introducción . . . . .                                                                                                | 105        |
| J.2. Método de Elemento Finito del Tipo Mínimos Cuadrados para Sistemas de Primer Orden . . . . .                            | 106        |
| J.2.1. Primer Sistema . . . . .                                                                                              | 106        |
| J.2.2. Segundo Sistema . . . . .                                                                                             | 108        |

# Índice de cuadros

|                                                              |    |
|--------------------------------------------------------------|----|
| 4.1. Orden de Convergencia $factor = 50$ (FDM) . . . . .     | 42 |
| 4.2. Orden de Convergencia $factor = 100$ (FDM) . . . . .    | 43 |
| 4.3. Orden de Convergencia $factor = 200$ (FDM) . . . . .    | 44 |
| 4.4. Orden de Convergencia $factor = 400$ (FDM) . . . . .    | 45 |
| 4.5. Orden de Convergencia $factor = 50$ (FEMGM) . . . . .   | 47 |
| 4.6. Orden de Convergencia $factor = 100$ (FEMGM) . . . . .  | 48 |
| 4.7. Orden de Convergencia $factor = 200$ (FEMGM) . . . . .  | 49 |
| 4.8. Orden de Convergencia $factor = 400$ (FEMGM) . . . . .  | 50 |
| 4.9. Orden de Convergencia $factor = 50$ (LSFEM) . . . . .   | 51 |
| 4.10. Orden de Convergencia $factor = 100$ (LSFEM) . . . . . | 52 |
| 4.11. Orden de Convergencia $factor = 200$ (LSFEM) . . . . . | 53 |
| 4.12. Orden de Convergencia $factor = 400$ (LSFEM) . . . . . | 54 |

|                                                                                                                      |    |
|----------------------------------------------------------------------------------------------------------------------|----|
| 4.13. $\tilde{k}$ para diversos valores de $h$ y $k$ para el Método de diferencias finitas implícito . . . . .       | 59 |
| 4.14. $\tilde{k}$ para diversos valores de $h$ y $k$ para el Método de elementos finitos Galerkin Mixto . . . . .    | 63 |
| 4.15. $\tilde{k}$ para diversos valores de $h$ y $k$ para el Método de elementos finitos Mínimos Cuadrados . . . . . | 67 |

# Índice de figuras

|      |                                                                   |    |
|------|-------------------------------------------------------------------|----|
| 2.1. | Funciones de forma lineales en el intervalo $[0, 1]$ . . . . .    | 16 |
| 2.2. | Funciones de forma cuadráticas en el intervalo $[0, 1]$ . . . . . | 18 |
| 2.3. | Funciones de forma de tercer grado por elemento . . . . .         | 19 |
| 2.4. | Polinomios de Hermite de tercer grado . . . . .                   | 21 |
| 4.1. | Error <i>factor</i> = 50 (FDM) . . . . .                          | 42 |
| 4.2. | Error <i>factor</i> = 100 (FDM) . . . . .                         | 43 |
| 4.3. | Error <i>factor</i> = 200 (FDM) . . . . .                         | 44 |
| 4.4. | Error <i>factor</i> = 400 (FDM) . . . . .                         | 45 |
| 4.5. | Error <i>factor</i> = 50 (FEMGM) . . . . .                        | 46 |
| 4.6. | Error <i>factor</i> = 100 (FEMGM) . . . . .                       | 47 |
| 4.7. | Error <i>factor</i> = 200 (FEMGM) . . . . .                       | 48 |
| 4.8. | Error <i>factor</i> = 400 (FEMGM) . . . . .                       | 49 |

|                                                                                                                                                                  |    |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 4.9. Error $factor = 50$ (LSFEM) . . . . .                                                                                                                       | 51 |
| 4.10. Error $factor = 100$ (LSFEM) . . . . .                                                                                                                     | 52 |
| 4.11. Error $factor = 200$ (LSFEM) . . . . .                                                                                                                     | 53 |
| 4.12. Error $factor = 400$ (LSFEM) . . . . .                                                                                                                     | 54 |
| 4.13. Solución exacta, interpolante y solución numérica . . . . .                                                                                                | 55 |
| 4.14. $\tilde{k}$ versus $k$ para diferentes valores de $h$ para el método de diferencias<br>finitas implícito . . . . .                                         | 61 |
| 4.15. Error porcentual en la aproximación de $k$ para el método de diferencias<br>finitas implícito . . . . .                                                    | 61 |
| 4.16. Gráficas de la parte real, imaginaria y valor absoluto de los diferentes<br>valores de $\lambda$ para el método de diferencias finitas implícito . . . . . | 62 |
| 4.17. $\tilde{k}$ versus $k$ para diferentes valores de $h$ para el método de elementos<br>finitos Galerkin Mixto . . . . .                                      | 64 |
| 4.18. Error porcentual en la aproximación de $k$ para el método de elementos<br>finitos Galerkin Mixto . . . . .                                                 | 65 |
| 4.19. Gráficas de la parte real, imaginaria y valor absoluto de los diferentes<br>valores de $\lambda$ para el método de elementos finitos Galerkin Mixto . . .  | 65 |
| 4.20. $\tilde{k}$ versus $k$ para diferentes valores de $h$ para el Método de elementos<br>finitos Mínimos Cuadrados . . . . .                                   | 67 |

|                                                                              |    |
|------------------------------------------------------------------------------|----|
| 4.21. Error porcentual en la aproximación de $k$ para el Método de elementos |    |
| finitos Mínimos Cuadrados . . . . .                                          | 68 |

# Introducción

Desde hace aproximadamente tres años un grupo de investigadores de la Universidad Central de Venezuela y la Universidad de Carabobo ha venido realizando un estudio del método de elementos finitos basados en mínimos cuadrados (LSFEM) sobre sistemas de ecuaciones diferenciales ordinarias y parciales de primer orden. En este trabajo definimos dicho método, presentamos las ventajas con respecto al tradicional método de elementos finitos Galerkin Mixto, y luego lo comparamos con algunos métodos en diferencias finitas. El estudio lo realizamos considerando la aplicación del método a un problema de dispersión modelado por la ecuación de Helmholtz en una dimension espacial (1-D). De acuerdo a nuestro conocimiento no existe en la literatura una comparación del método de LSFEM aplicado a problemas exteriores descritos por la ecuación de Helmholtz, con los otros métodos estudiados en este trabajo. Se hará énfasis en el orden de convergencia, así como también en un estudio del error debido a la dispersión para los métodos de elementos finitos Galerkin Mixto y Mínimos Cuadrados para sistemas de primer orden.

La organización de este documento se hace en base a capítulos. En el primer capítulo se hace el planteamiento del problema a resolver, se presentan además los objetivos de esta investigación, un breve enfoque metodológico, así como la so-

lución analítica y la transformación de la ecuación de Helmholtz en un sistema de ecuaciones diferenciales de primer orden. En el segundo capítulo se hará una introducción a los métodos numéricos a comparar en este trabajo, comenzando con el método de diferencias finitas en el que se mencionan un método implícito. Luego se desarrollan de forma general los métodos de elementos finitos Galerkin Mixto y Mínimos Cuadrados. En el tercer capítulo se aplican los métodos descritos en el capítulo dos para resolver el problema de dispersión de ondas unidimensional y en el cuarto capítulo se presentan los resultados basados en el análisis de orden de convergencia de los métodos así como en el análisis de la contaminación del error debido en la dispersión y por último las conclusiones y trabajos futuros. Además de esto en los apendices se presentan los programas desarrollados en Matlab y Maple; también se presentan avances realizados con el método de diferencias finitas dominio tiempo así como para la aplicación del método de elementos finitos mínimos cuadrados a un problema similar en dos dimensiones.

# Capítulo 1

## PLANTEAMIENTO DEL PROBLEMA

### 1.1. Formulación del Problema de Dispersión de Ondas en 1-D

Desde el punto de vista de las aplicaciones, la ecuación de Helmholtz modela fenómenos de propagación de ondas en la naturaleza, en particular problemas de dispersión, los cuales son problemas exteriores definidos sobre dominios no acotados. La ecuación de onda reducida o ecuación de Helmholtz en dos o tres dimensiones espaciales está dada por

$$\Delta p + \frac{\omega^2}{c^2} p = 0 \tag{1.1}$$

donde  $\omega$  es la frecuencia temporal,  $c$  es la velocidad de la onda, y  $p$  representa la presión. La relación  $k \equiv \frac{\omega}{c}$  es denominada “numero de onda”. En este trabajo consideraremos el problema de dispersión de una onda incidente,  $p_{inc} = e^{ikx}e^{-i\omega t}$ , a partir de una pared que se extiende infinitamente por todos sus lados (problema

unidimensional, 1D). La ecuación que modela este fenómeno es

$$p'' + k^2 p = 0, \quad 0 < x < \infty \quad (1.2)$$

sujeta a la condición de frontera de pared rígida  $p'(0) = 0$ . Para que este problema resulte ser un problema bien planteado es necesario agregar una condición de irradiación en el borde infinito. Esta condición fue introducida por Sommerfeld y en este caso equivale a la ecuación satisfecha por la onda dispersada que viaja hacia la derecha en el borde infinito.

$$\frac{dp_{sc}}{dx} - ikp_{sc} = 0 \quad x \rightarrow \infty \quad (1.3)$$

En el caso unidimensional se tiene la particularidad de que la condición de irradiación de Sommerfeld (1.3) se satisface exactamente en cada punto del dominio,  $x > 0$ . La presión,  $p(x)$ , puede ser descompuesta en la suma de la presión para onda incidente y la de la onda esparcida,  $p(x) = p_{inc}(x) + p_{sc}(x)$ . Luego el problema a resolver para la onda esparcida  $p_{sc}$  sería

$$p_{sc}'' + k^2 p_{sc} = 0, \quad 0 < x < 1 \quad (1.4)$$

$$p_{sc}'(0) = ik \quad (1.5)$$

$$\frac{dp_{sc}}{dx}(1) - ikp_{sc}(1) = 0 \quad (1.6)$$

A partir de dicha solución y la presión para la onda incidente se puede obtener la presión total  $p(x)$ ,

## 1.2. Objetivo General

El objetivo principal de este trabajo es resolver numéricamente la ecuación de Helmholtz en una dimensión por el método de elementos finitos basado en míni-

mos cuadrados y compararlo con otros métodos numéricos para identificar las fortalezas y debilidades del método.

### 1.3. Objetivos Específicos

- (a) Formulación del método de los elementos finitos basados en mínimos cuadrados.
- (b) Aplicación a problemas de dispersión de ondas en una dimensión espacial.
- (c) Formulación del método de elementos finitos del tipo Galerkin mixto para la ecuación de Helmholtz en una dimensión.
- (d) Formulación de un método implícito en diferencias finitas para el mismo problema de dispersión.
- (e) Estudiar los fenómenos de contaminación y dispersión relacionados con la solución numérica de la ecuación de Helmholtz por el método de elementos finitos.
- (f) Desarrollar en *MATLAB<sup>TM</sup>* códigos correspondientes a los métodos mencionados en los puntos a., c., d. y e.
- (g) Comparar los resultados obtenidos y sacar conclusiones en cuanto al orden de convergencia y la contaminación del error debido a la dispersión.

### 1.4. Enfoque Metodológico

El problema dado por (1.4), (1.5) y (1.6) será transformado en un sistema de ecuaciones diferenciales de primer orden y luego resuelto numéricamente por los

métodos de diferencias finitas, elementos finitos Galerkin mixto y elementos finitos del tipo mínimos cuadrados para un sistema de ecuaciones diferenciales de primer orden. El método de diferencias finitas se basa en la discretización del dominio y posterior sustitución de la aproximación de las derivadas en las formas de diferencias, las cuales se obtienen bien a partir de una Serie de Taylor o por medio de la diferenciación de un polinomio de interpolación de Lagrange; de esta manera se pueden obtener esquemas explícitos, implícitos o semi-implícitos, en nuestro caso se generará un esquema implícito que conllevará a un sistema de ecuaciones lineales, a resolver. Los métodos de elementos finitos también hacen uso de la discretización del dominio para luego utilizar una forma residual pesada. Dependiendo de la selección de dicho peso tendremos el método de Galerkin o el de mínimos cuadrados, entre otros. En ambos casos, se sustituyen funciones de aproximación y de prueba que nuevamente nos permiten obtener un sistema de ecuaciones lineales. Por último, se resuelve dicho sistema para obtener así una aproximación a la solución deseada.

## 1.5. Área de Conocimiento

Análisis Numérico, en particular Solución Numérica de Ecuaciones Diferenciales por los métodos de diferencias finitas y elementos finitos para el estudio de la propagación de ondas en medios acústicos.

## 1.6. Solución Analítica

Para resolver la ecuación diferencial ordinaria lineal con coeficientes constantes (1.4) sujeta a las condiciones de frontera (1.5) y (1.6), se utilizará el procedimiento básico para este tipo de ecuación, ver [Kis].

Sea la ecuación característica  $\lambda^2 + k^2 = 0$  que tiene por solución  $\lambda = \pm ik$ , entonces el sistema fundamental de la ecuación (1.4) es de la forma  $\cos(kx)$  y  $\sin(kx)$  y su solución general  $p_{sc}(x) = A \cos(kx) + B \sin(kx)$ .

De las condiciones (1.5) y (1.6) se obtiene que  $B = i$  y  $A = 1$ , por lo que la solución general sería:

$$p_{sc}(x) = \cos(kx) + i \sin(kx) = e^{ikx} \quad (1.7)$$

Ésta solución es la que en el capítulo cuatro se comparará con la solución numérica.

## 1.7. Sistema de Ecuaciones Diferenciales de Primer Orden

Para aplicar los métodos numéricos en la resolución de la ecuación de Helmholtz, tal cual está planteado en este trabajo, es necesario conformar un sistema de ecuaciones diferenciales ordinarias de primer orden en una variable, el cual puede ser escrito en la forma matricial:

$$A_1 \frac{du}{dx} + A_0 u = f \quad (1.8)$$

donde  $u^t = (u_1, u_2, \dots, u_m)$  es un vector de  $m$  funciones desconocidas de  $x$ ,  $A_1 = (a_{ij,1})$ ,  $A_0 = (a_{ij,0})$  son matrices de orden  $m$  y  $f$  es un vector de funciones conocidas que dependen de  $x$ .

Para transformar el problema (1.4), (1.5) y (1.6), el cual es objeto de estudio, en un sistema de primer orden como el dado por (1.8) se utiliza el cambio de variables  $z = p'$ , obteniéndose el sistema con condiciones de frontera dado por:

$$z - p' = 0 \tag{1.9}$$

$$z' + k^2 p = 0 \tag{1.10}$$

$$z(0) - ik = 0 \tag{1.11}$$

$$z(1) - ikp(1) = 0 \tag{1.12}$$

El sistema dado por las ecuaciones (1.9) y (1.10) se puede expresar en la forma matricial  $Au = (A_0 + A_1 \frac{d}{dx})u = 0$ , donde  $A_0 = \begin{bmatrix} 1 & 0 \\ 0 & k^2 \end{bmatrix}$ ,  $A_1 = \begin{bmatrix} 0 & -1 \\ 1 & 0 \end{bmatrix}$ ,  $u = \begin{bmatrix} z \\ p \end{bmatrix}$

y  $0 = \begin{bmatrix} 0 \\ 0 \end{bmatrix}$ . A menos que se establezca lo contrario, esta forma matricial es la que se empleará para resolver el problema de dispersión de ondas en 1-D.

## Capítulo 2

# FORMULACIÓN GENERAL DE LOS MÉTODOS NUMÉRICOS

Diversos métodos numéricos han sido desarrollados para resolver ecuaciones diferenciales de manera aproximada. En el siglo veinte, hubo un crecimiento sin precedentes de dichos métodos, creando muchas áreas de investigación . En el *Journal of Computational and Applied Mathematics*, 128 (2001) se presentan varias contribuciones por las más famosas autoridades científicas en cada una de dichas áreas. El primer artículo [Tho] describe la historia del análisis numérico de ecuaciones diferenciales parciales, comenzando por el desarrollo de los métodos de diferencias finitas para problemas elípticos desde la década de los años 30 y el estudio intenso de dichos métodos para problemas de valor inicial durante el período 1950-1970, también se hace referencia a los métodos de elementos finitos desde el punto de vista de la ingeniería estructural, así como al subsecuente desarrollo desde el punto de vista del análisis matemático con funciones de aproximación polinomial por intervalos y a diversas clases de métodos tales como colocación, espectrales, de volumen finito y elementos de frontera. En los próximos tres artículos [Bia], [Got] y [Dah] se describen los métodos de colocación

con Splines, espectrales y ondículas. En otro artículo [Ran] se desarrollan métodos adaptativos, así como métodos de elementos finitos Galerkin discontinuo [Coc]. Otro [Sur], hace énfasis a las versiones  $p$  y  $hp$  de los métodos de elementos finitos. Esto es solo una minúscula muestra de la investigación desarrollada en el siglo veinte alrededor de la solución numérica de ecuaciones en derivadas parciales, entre los que se encuentran los que se tratan en este trabajo. En la próxima sección se dará una introducción genérica a cada uno de ellos.

## 2.1. Introducción del Método de Diferencias Finitas

La idea básica del método de diferencias finitas es reemplazar las derivadas por diferencias finitas. En este documento se presentará un enfoque basado en la serie de Taylor en una variable centrada en  $x = x_0$

$$f(x) = f(x_0) + f'(x_0)(x - x_0) + \frac{1}{2}f''(x_0)(x - x_0)^2 + \cdots + \frac{1}{n!}f^{(n)}(x_0)(x - x_0)^n + \cdots \quad (2.1)$$

Al sustituir en (2.1)  $x = x_0 + h$ ,  $h > 0$ , se obtiene

$$f(x_0 + h) = f(x_0) + f'(x_0)h + \frac{1}{2}f''(x_0)h^2 + \cdots + \frac{1}{n!}f^{(n)}(x_0)h^n + \cdots \quad (2.2)$$

que al despejar la primera derivada de  $f$  evaluada en  $x_0$  se obtiene la comúnmente denominada fórmula de diferencias hacia adelante

$$f'(x_0) = \frac{f(x_0 + h) - f(x_0)}{h} - \frac{1}{2}f''(x_0)h + \cdots - \frac{1}{n!}f^{(n)}(x_0)h^{n-1} + \cdots \quad (2.3)$$

Si en esta última ecuación se sustituye  $h$  por  $-h$  se obtiene la fórmula de diferencias hacia atrás

$$f'(x_0) = \frac{f(x_0) - f(x_0 - h)}{h} + \frac{1}{2}f''(x_0)h + \cdots + \frac{(-1)^n}{n!}f^{(n)}(x_0)h^{n-1} + \cdots \quad (2.4)$$

Al aproximar  $f'(x_0)$  por el primer término de (2.3) o (2.4) se obtiene un error de orden uno, el que expresamos como  $O(h)$ . Si se desea obtener una fórmula de aproximación de orden dos,  $O(h^2)$ , se puede hacer por el siguiente procedimiento:

- (1) En la ecuación (2.1) sustituir  $x = x_0 - h$  y obtener

$$f(x_0 - h) = f(x_0) - f'(x_0)h + \frac{1}{2}f''(x_0)h^2 + \cdots + \frac{(-1)^n}{n!}f^{(n)}(x_0)h^n + \cdots \quad (2.5)$$

- (2) Realizar la resta de las ecuaciones (2.2) y (2.5) con la finalidad de eliminar los términos que involucren derivadas de orden par, obteniéndose

$$\begin{aligned} f(x_0 + h) - f(x_0 - h) &= 2f'(x_0)h + \frac{2}{3!}f'''(x_0)h^3 + \cdots \\ &\quad + \frac{2}{(2n+1)!}f^{(2n+1)}(x_0)h^{2n+1} + \cdots \end{aligned} \quad (2.6)$$

- (3) Despejar  $f'(x_0)$ .

$$\begin{aligned} f'(x_0) &= \frac{f(x_0 + h) - f(x_0 - h)}{2h} - \frac{1}{3!}f'''(x_0)h^2 + \cdots \\ &\quad - \frac{1}{(2n+1)!}f^{(2n+1)}(x_0)h^{2n} + \cdots \end{aligned} \quad (2.7)$$

Para obtener otras aproximaciones para  $f'(x_0)$ ,  $f''(x_0)$ ,  $\cdots$ ,  $f^n(x_0)$  se pueden establecer procedimientos similares que involucren mayor cantidad de puntos. Existen muchos métodos para generar ecuaciones en diferencia para aproximar derivadas, entre los más utilizados que no se han mencionado aún se encuentran la diferenciación de polinomios de interpolación de Lagrange y un método de coeficientes indeterminados, en cualquier caso se puede consultar [Str], [Eva] y [Bur].

## 2.2. Formulación del Método de Diferencias Finitas Implícito

Al sustituir como se mencionó anteriormente la derivada por su aproximación en base a diferencias finitas se pueden generar los denominados esquemas explícitos, implícitos y semi-implícitos. En esta sección, se presentará de forma general y brevemente una idea de los esquemas de diferencias finitas implícitos. Para ello consideraremos la ecuación diferencial  $Lu = f$  en  $\Omega$ , donde  $L$  es un operador diferencial,  $u$  sería la función incógnita de la ecuación diferencial,  $f$  una función (en principio continua en  $\Omega$ ) y  $\Omega$  un dominio acotado en una, dos o tres dimensiones. Al utilizar aproximaciones en forma de diferencias finitas como las presentadas anteriormente u otras según las necesidades del operador diferencial  $L$  se obtiene un operador diferencial discreto  $L_d$  en diferencias finitas que al aplicarlo a  $u$  se obtiene una ecuación en diferencias  $L_d u_i = f_i$ , donde  $u_i$  y  $f_i$  son las funciones  $u$  y  $f$  evaluadas en un punto  $x_i$  del dominio, respectivamente. De esta forma al discretizar el dominio y utilizar tantas ecuaciones en diferencias como nodos involucrados en la discretización, considerando las condiciones de frontera, se obtiene un sistema de ecuaciones que se debe resolver usualmente por métodos iterativos ya que el sistema obtenido generalmente es esparcido. Al resolver dicho sistema de ecuaciones se obtienen aproximaciones de la función  $u$  en los nodos  $x_i$ .

## 2.3. Introducción del Método de Elementos Finitos

El método de los elementos finitos es una técnica computacional para obtener soluciones aproximadas a las ecuaciones diferenciales. Más que aproximar direc-

tamente la solución de una ecuación diferencial el método de elementos finitos utiliza un problema variacional que involucra una integral de la ecuación diferencial, sobre el dominio del problema. Este dominio es dividido en subdominios que son denominados elementos finitos y la solución de la ecuación diferencial es aproximada por una función polinomial sobre cada elemento. Una vez que ha sido hecho esto, la integral variacional es evaluada como la suma de contribuciones de cada elemento finito. De ello resulta un sistema algebraico que al resolverlo se obtienen aproximaciones a la solución de la ecuación diferencial. Para describir los diversos métodos de elementos finitos se hará el enfoque basado en los métodos de residuos ponderados, para ello primero se presentan algunas consideraciones comunes a dichos métodos.

### 2.3.1. Consideraciones Comunes

Sea la ecuación diferencial  $Lu = f$  dada en un dominio  $\Omega$ , donde  $L$  es un operador diferencial y  $f$  es una función suficientemente suave; con condición de frontera  $Bu = g$  sobre la frontera  $\gamma$ , definimos los residuos como  $R_\Omega = Lu_N - f$  y  $R_\gamma = Bu_N - g$  definidos sobre el dominio y la frontera respectivamente, donde  $u_N$  es una aproximación a la solución de la ecuación diferencial, la cual expresaremos como una combinación lineal de los elementos de la base  $\{\phi_1, \phi_2, \dots, \phi_i, \dots, \phi_N\}$  el cual es un espacio de funciones de dimensión finita, o sea

$$u_N(x) = \sum_{i=1}^N C_i \phi_i(x) \quad (2.8)$$

a las funciones  $\phi_i$  se le denominan usualmente funciones de forma. Sean también los productos internos  $(R_\Omega, v)$  y  $(R_\gamma, w)$  definidos por

$$(R_\Omega, v) = \int_{\Omega} v R_\Omega d\Omega \quad (2.9)$$

y

$$(R_\gamma, w) = \oint_\gamma w R_\gamma d\gamma \quad (2.10)$$

La forma integral equivalente se obtiene sumando los productos internos definidos con anticipación e igualando a cero, como se presenta a continuación

$$(R_\Omega, v) + (R_\gamma, w) = 0 \quad (2.11)$$

Al sustituir en la ecuación anterior la expresión (2.8) en conjunto con  $N$  pares de funciones de pesos  $v$  y  $w$  se obtiene un sistema de ecuaciones lineales de orden  $N$  que al resolverlo se tienen los valores de  $C_i$ ,  $i = 1, \dots, N$ . Para más detalle se puede consultar [Red], [Eva], [Joh] y [Fla]. Dependiendo de la selección de la funciones de peso  $v$  y  $w$  se tienen diversos métodos, algunos de los cuales se presentan a continuación.

### **Método de Colocación por Puntos**

Este método se basa en fijar  $N$  puntos  $x_i$  diferentes del dominio  $\Omega$  y seleccionar como funciones de peso la función *delta de Dirac*  $\delta(x - x_i)$ ,  $i = 1, \dots, N$ . Al hacer esto se genera, como se mencionó anteriormente, un sistema de  $N$  ecuaciones con  $N$  incógnitas, cuya solución  $C_i$  al sustituirla en (2.8) permite obtener una aproximación a la solución  $u$  de la ecuación diferencial.

## Método de Colocación por Subdominios

Este método consiste en dividir el dominio  $\Omega$  en  $N$  subdominios,  $\Omega_i$ ,  $i = 1, \dots, N$  y seleccionar funciones de peso, según la expresión

$$v(x) = w(x) = \begin{cases} 1 & \text{si } x \in \Omega_i, \\ 0 & \text{en otros casos.} \end{cases} \quad (2.12)$$

## Método de Galerkin

Este es un método muy popular, el cual consiste en seleccionar las funciones de peso iguales a las funciones de forma, es decir  $v_i = w_i \equiv \phi_i$ . Usualmente se seleccionan funciones de forma tal que satisfagan las condiciones de frontera, en tal caso el producto interno dado por (2.10) es cero.

## Método de Mínimos Cuadrados

Para este método se define la integral

$$I = (R_\Omega, R_\Omega) + (R_\gamma, R_\gamma) \quad (2.13)$$

donde los productos internos se definen según (2.9) y (2.10), así el problema de resolver (2.11) sería equivalente a minimizar (2.13), para lo cual aplicamos la condición necesaria  $\delta I = 0$ , obteniendo el sistema de ecuaciones dado por

$$\frac{\partial}{\partial C_j} [(R_\Omega, R_\Omega) + (R_\gamma, R_\gamma)] = 0 \quad j = 1, \dots, N \quad (2.14)$$

Recuérdese que los residuos  $R_\Omega$  y  $R_\gamma$  depende de  $u_N$  según (2.9) y (2.10), y ésta a su vez depende de  $C_j$  (2.8)

## 2.4. Método de Elementos Finitos

El método de los elementos finitos establece un procedimiento sistemático para generar funciones de forma sobre el dominio en el que se desee resolver la ecuación diferencial, bien sea ordinaria o parcial. A continuación, se presenta el procedimiento usualmente utilizado para el método de elementos finitos

- a. Dividir el dominio en elementos finitos que no son más que subdominios geoméricamente simples, cuya unión sea el dominio.
- b. Construir las funciones de forma. Para esto se utilizan usualmente polinomios de interpolación de Lagrange que cumplan la condición

$$\phi_j(x_i) = \delta_{ij} \quad (2.15)$$

donde  $\delta_{ij}$  es la delta de Kronecker, siendo  $i$  el nodo del elemento donde  $\phi_i$  vale uno y  $j$  cualquiera de los nodos pertenecientes a los elementos en que fue dividido el dominio.

- c. Sustituir las funciones de forma en la expresión (2.8). De esta manera los  $C_i$  corresponderían a los valores de  $u$  en el nodo  $i$  a los cuales denotaremos por  $U_i$ , teniendo así que

$$u_N = \sum_{i=1}^N U_i \phi_i \quad (2.16)$$

- d. Sustituir esta última expresión en la ecuación que rige el método de los residuos ponderados y seleccionar entre los métodos de colocación puntual, colocación por subdominios, Galerkin, Mínimos Cuadrados, entre otros.

e. Resolver el sistema de ecuaciones lineales donde las incógnitas serían los  $U_i$ ,  $i = 1, \dots, N$

De esta forma utilizando (2.8) se tiene una aproximación a la solución de la ecuación diferencial.

### 2.4.1. Construcción de funciones de Forma

Debido a que en este trabajo se resuelve el problema de Dispersión de Onda Unidimensional es necesario describir en detalle la construcción de funciones de forma a trozos sobre un intervalo de  $R$ , sin pérdida de generalidad se hará sobre una discretización uniforme del intervalo  $[0, 1]$ . A cada subintervalo de la discretización se le denomina elemento. Un elemento puede contener más de dos nodos, ello va a depender del grado del polinomio de interpolación a utilizar. La cantidad de nodos es igual al grado del polinomio mas uno, esto es debido a que se desea que el polinomio de interpolación evaluado en los nodos sea cero excepto en un punto en donde debe valer uno.

#### Funciones Lineales

Para describir las funciones de forma lineales se dividirá el intervalo  $[0, 1]$  en cuatro intervalos de igual longitud, generándose así cinco nodos, dos por cada elemento. Se procederá a detallar la función de forma que vale uno en el tercer nodo y cero en los demás. Por ello, esta función de forma es cero en el primer y cuarto elemento, y en los elementos dos y tres son funciones lineales creciente y decreciente respectivamente. De igual manera se puede analizar las otras cuatro

funciones de forma. Observe la función de forma de color rojo en la figura 2.1. En general se pueden obtener funciones de forma lineales dividiendo el intervalo

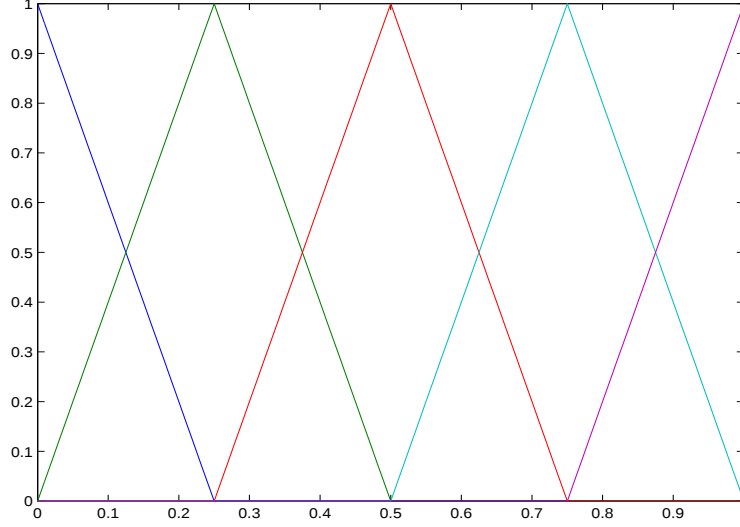


Figura 2.1: Funciones de forma lineales en el intervalo  $[0, 1]$

$[0, 1]$  en  $N$  subintervalos denotando por  $x_0, x_1, \dots, x_{i-1}, x_i, x_{i+1}, \dots, x_N$  los nodos ordenados de menor a mayor, definiendo la función  $\phi_i$  por medio de la expresión

$$\phi_i(x) = \begin{cases} \frac{x-x_{i-1}}{x_i-x_{i-1}} & \text{si } x_{i-1} < x \leq x_i \\ \frac{x-x_{i+1}}{x_i-x_{i+1}} & \text{si } x_i < x < x_{i+1}, \\ 0 & \text{en otros casos.} \end{cases} \quad (2.17)$$

se puede observar que ella cumplen con las condiciones mencionadas anteriormente, o sea  $\phi_i(x_i) = 1$  y  $\phi_i(x_j) = 0$ ,  $j \neq i$ .

## Funciones de Grado Superior

Un criterio similar se utiliza para generar funciones de forma de segundo y tercer grado, la diferencia se centra en que una vez dividido el intervalo en elementos

cada uno de ellos tendrán tres y cuatro nodos respectivamente. En cada elemento, la función de forma debe ser uno en un solo nodo y cero en los demás. Observe las figuras 2.2. y 2.3.

En la figura 2.2 se presentan funciones de forma cuadráticas, en este caso se divide el intervalo  $[0, 1]$  en tres intervalos (*elementos*) de igual longitud, por cada elemento se consideran tres nodos (*los de los extremos y el punto medio*), obteniéndose de esta manera siete funciones de forma (*cada una con un color diferente*), cuyas ecuaciones son

$$\phi_1(x) = \begin{cases} 18(\frac{1}{3} - x)(\frac{1}{6} - x) & \text{si } 0 \leq x < \frac{1}{3}, \\ 0 & \text{en otros casos.} \end{cases} \quad (2.18)$$

$$\phi_2(x) = \begin{cases} 36x(\frac{1}{3} - x) & \text{si } 0 \leq x < \frac{1}{3}, \\ 0 & \text{en otros casos.} \end{cases} \quad (2.19)$$

$$\phi_3(x) = \begin{cases} -18x(\frac{1}{6} - x) & \text{si } 0 \leq x < \frac{1}{3}, \\ 18(\frac{1}{2} - x)(\frac{2}{3} - x) & \text{si } \frac{1}{3} \leq x < \frac{2}{3}, \\ 0 & \text{en otros casos.} \end{cases} \quad (2.20)$$

$$\phi_4(x) = \begin{cases} -36(\frac{2}{3} - x)(\frac{1}{3} - x) & \text{si } \frac{1}{3} \leq x < \frac{2}{3}, \\ 0 & \text{en otros casos.} \end{cases} \quad (2.21)$$

$$\phi_5(x) = \begin{cases} 18(\frac{1}{3} - x)(\frac{1}{2} - x) & \text{si } \frac{1}{3} \leq x < \frac{2}{3}, \\ 18(1 - x)(\frac{5}{6} - x) & \text{si } \frac{2}{3} \leq x < 1, \\ 0 & \text{en otros casos.} \end{cases} \quad (2.22)$$

$$\phi_6(x) = \begin{cases} -36(\frac{2}{3} - x)(1 - x) & \text{si } \frac{2}{3} \leq x < 1, \\ 0 & \text{en otros casos.} \end{cases} \quad (2.23)$$

$$\phi_7(x) = \begin{cases} 18(\frac{2}{3} - x)(\frac{5}{6} - x) & \text{si } \frac{2}{3} \leq x \leq 1, \\ 0 & \text{en otros casos.} \end{cases} \quad (2.24)$$

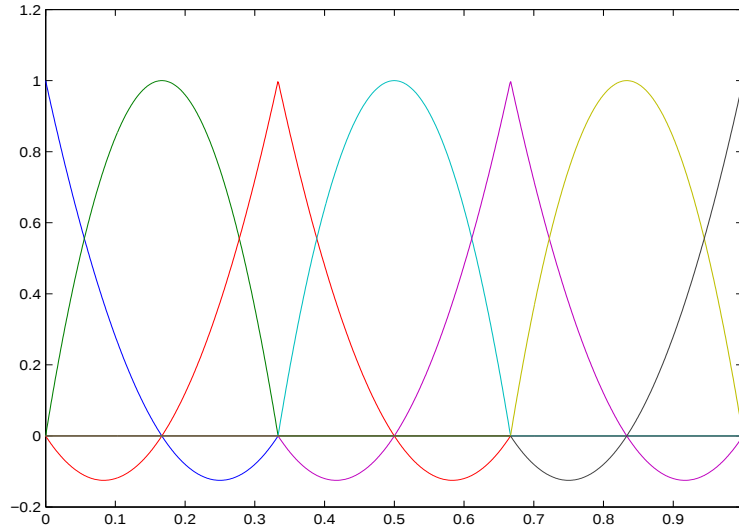


Figura 2.2: Funciones de forma cuadráticas en el intervalo  $[0, 1]$

Para estudiar las funciones cúbicas vamos a ubicarnos en otro contexto. Asumamos que se ha realizado una transformación de cada elemento (por medio de un cambio de variables) y así todos los elementos se transforma en uno equivalente en el intervalo  $[-1, 1]$ . Luego se ubican cuatro nodos, dos en los extremos y

los otros en el interior del intervalo en cuestión, tal que los nodos estén espaciados uniformemente. Se calculan las funciones de forma que cumplan la condición (2.15) obteniéndose los polinomios mostrados en la figura 2.3, los cuales tienen por ecuación

$$\phi_1(x) = -\frac{9}{16}(x-1)(x-\frac{1}{3})(x+\frac{1}{3}) \quad (2.25)$$

$$\phi_2(x) = \frac{27}{16}(x+1)(x-\frac{1}{3})(x-1) \quad (2.26)$$

$$\phi_3(x) = -\frac{27}{16}(x-1)(x+1)(x+\frac{1}{3}) \quad (2.27)$$

$$\phi_4(x) = \frac{9}{16}(x+1)(x-\frac{1}{3})(x+\frac{1}{3}) \quad (2.28)$$

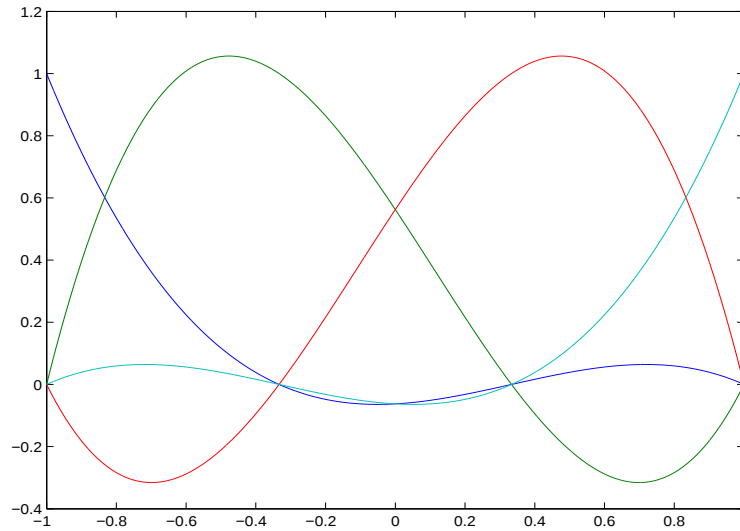


Figura 2.3: Funciones de forma de tercer grado por elemento

## Polinomios de Hermite

La limitación básica de las funciones de forma presentadas anteriormente es la no continuidad de las derivadas en los nodos extremos de los elementos. Para ello

se pueden seleccionar funciones de forma de otro espacio de funciones. Algunas funciones utilizadas con frecuencia involucran Splines, Polinomios de Hermite, entre otras. Aquí se presenta una introducción a los Polinomios de Hermite, los cuales se definen en el intervalo  $[-1,1]$  de tal manera que la función y su derivada evaluada en los extremos del intervalo sean cero excepto una sola de ellas, en cuyo caso tomará el valor de uno. Es de destacar, que hay que hacer una transformación del elemento de dos nodos al intervalo  $[-1,1]$ .

Comencemos con un polinomio de tercer grado. Se desea que el polinomio cumpla las siguientes condiciones  $P(-1) = 1$ ,  $P(1) = 0$ ,  $\frac{dP}{dx}(-1) = 0$  y  $\frac{dP}{dx}(1) = 0$  en cuyo caso se obtiene el polinomio de tercer grado

$$\phi_1(x) = \frac{1}{4}(1-x)^2(x+2) \quad (2.29)$$

De igual manera se calculan los otros tres polinomios de Hermite de grado tres, obteniéndose las ecuaciones

$$\phi_2(x) = \frac{1}{4}(1+x)^2(2-x) \quad (2.30)$$

$$\phi_3(x) = \frac{1}{4}(1-x)^2(x+1) \quad (2.31)$$

$$\phi_4(x) = \frac{1}{4}(1+x)^2(x-1) \quad (2.32)$$

Observe dichos polinomios en la figura 2.4

Al sustituir los polinomios de Hermite en (2.16) se tiene que

$$u_N = \sum_{i=1}^2 (U_i \phi_i(x) + DU_i \phi_{i+2}(x)) \quad (2.33)$$

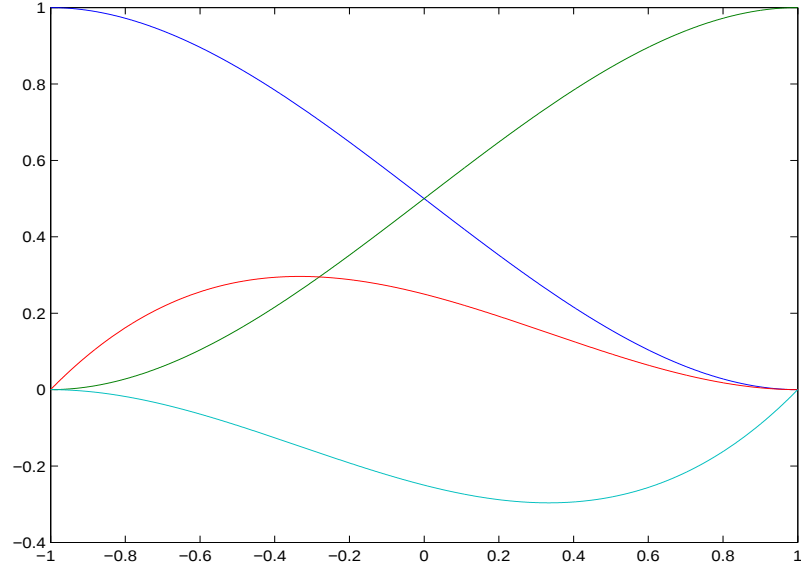


Figura 2.4: Polinomios de Hermite de tercer grado

donde  $U_i$  al igual que en los casos anteriores representa el valor aproximado de  $u$  en el nodo  $i$  y  $DU_i$  la aproximación de la derivada de  $u$  evaluada en el mismo nodo. Fácilmente se puede demostrar que  $u_N$  es una función continuamente diferenciable en el dominio. Las funciones  $\phi_3(x)$  y  $\phi_4(x)$  son tales que al evaluarles la derivada en  $x = -1$  y  $x = 1$  se obtiene el valor uno. De una manera similar se puede proceder para obtener polinomios de Hermite de grado 5, 7, 9, etc., en cuyo caso se necesitan las segundas, terceras y cuartas derivadas evaluadas en los extremos.

### 2.4.2. Formulación del Método de Elementos Finitos Galerkin Mixto

Para hacer una descripción más o menos genérica del método de elementos finitos Galerkin Mixto, se partirá de la ecuación diferencial ordinaria lineal de orden  $n$  con coeficientes constantes

$$a_n u_1^{(n)} + a_{n-1} u_1^{(n-1)} + \cdots + a_2 u_1'' + a_1 u_1' + a_0 u_1 = f_1 \quad (2.34)$$

Es bien sabido [Guz] que haciendo los cambios de variables adecuados se puede transformar esta ecuación en el siguiente sistema de ecuaciones diferenciales de primer orden

$$u_1' = u_2 \quad (2.35)$$

$$u_2' = u_3 \quad (2.36)$$

$$\vdots$$

$$u_{n-3}' = u_{n-2} \quad (2.37)$$

$$u_{n-2}' = u_{n-1} \quad (2.38)$$

$$u_{n-1}' = u_n \quad (2.39)$$

$$a_n u_n' + a_{n-1} u_{n-1}' + \cdots + a_2 u_2' + a_1 u_1' + a_0 u_1 = f_1 \quad (2.40)$$

El sistema (2.35-2.40) se puede expresar en forma matricial como  $Au = f$ , donde

$$A = \begin{bmatrix} a_0 & 0 \\ 0 & I \end{bmatrix} + \begin{bmatrix} b & a_n \\ -I & 0 \end{bmatrix} \frac{d}{dx}, \quad u = \begin{bmatrix} u_1 & u_2 & \cdots & u_n \end{bmatrix}^t \text{ y } f = \begin{bmatrix} f_1 & 0 & \cdots & 0 & 0 \end{bmatrix}^t,$$

siendo  $b = \begin{bmatrix} a_1 & a_2 & a_3 & \cdots & a_{n-1} \end{bmatrix}$  e  $I$  la matriz identidad de orden  $n - 1$

Se define luego el residuo  $R_\Omega = Au_N - f$ , donde  $u_N$  es una aproximación al vector solución  $u$  del sistema de ecuaciones diferenciales de primer orden (2.35-2.40), expresando cada elemento de dicho vector como una combinación lineal de vectores, cuyas componentes son funciones de forma seleccionadas con anticipación, tal como en (2.8). Se define también el producto interno, similar a (2.9), como

$$(R_\Omega, v) = \int_\Omega \bar{v}^t R_\Omega d\Omega \quad (2.41)$$

donde  $\bar{v}^t = \begin{bmatrix} \bar{v}_1 & \bar{v}_2 & \cdots & \bar{v}_n \end{bmatrix}$  es el vector de la conjugada compleja de las funciones de peso.

La forma integral equivalente sería

$$(R_\Omega, v) = 0 \quad (2.42)$$

Al seleccionar las funciones de forma y sustituirlas en (2.42), integrar por partes los términos que permitan incorporar las condiciones de frontera para el problema dado por (2.35-2.40) se genera un sistema de ecuaciones lineales, que al resolver se obtiene la aproximación deseada, como lo establece en las consideraciones comunes del método de los elementos finitos, ya presentado.

### 2.4.3. Formulación del Método de Elementos Finitos Mínimos Cuadrados para Sistemas de Primer Orden

Dada la ecuación diferencial lineal  $Au = f$ , sujeta a las condiciones de frontera  $Bu = g$ , donde  $A$  y  $B$  son operadores diferenciales lineales. Se desea minimizar el funcional  $I(u) = \frac{1}{2}\|Au - f\|_0^2 + \frac{1}{2}\|Bu - g\|_\gamma^2$  para ello se sabe, [Fox] y [Gel], que una condición necesaria es

$$\lim_{t \rightarrow 0} \frac{dI(u + tv)}{dt} = 0 \quad (2.43)$$

Tenemos que

$$I(u + tv) = \frac{1}{2}(A(u + tv) - f, A(u + tv) - f) + \frac{1}{2}\langle B(u + tv) - g, B(u + tv) - g \rangle_\gamma \quad (2.44)$$

aplicando las propiedades del producto interno y luego derivando respecto a  $t$  se obtiene

$$\begin{aligned} \frac{dI(u + tv)}{dt} &= \frac{1}{2}[(Au, Av) + (Av, Au) + 2t(Av, Av) - (f, Av) - (Av, f)] \\ &\quad + \frac{1}{2}[\langle Bu, Bv \rangle_\gamma + \langle Bv, Bu \rangle_\gamma + 2t\langle Bv, Bv \rangle_\gamma - \langle g, Bv \rangle_\gamma - \langle Bv, g \rangle_\gamma] \end{aligned} \quad (2.45)$$

luego al aplicar el límite cuando  $t$  tiende a cero e igualando a cero

$$\begin{aligned} \lim_{t \rightarrow 0} \frac{dI(u + tv)}{dt} &= \frac{1}{2}[(Au, Av) + (Av, Au) - (f, Av) - (Av, f)] \\ &+ \frac{1}{2}[\langle Bu, Bv \rangle_\gamma + \langle Bv, Bu \rangle_\gamma - \langle g, Bv \rangle_\gamma - \langle Bv, g \rangle_\gamma] = 0 \end{aligned} \quad (2.46)$$

Basándonos en uno de los axiomas para el producto interno de un espacio vectorial complejo, el cual establece que  $(x, y) = \overline{(y, x)}$  y como  $(x, y) + \overline{(x, y)} = 2\text{Re}(x, y)$  se concluye de (2.46) que

$$\text{Re}(Au, Av) + \text{Re}\langle Bu, Bv \rangle_\gamma = \text{Re}(f, Av) + \text{Re}\langle g, Bv \rangle_\gamma \quad (2.47)$$

Al sustituir en la ecuación anterior la expresión (2.8) se obtiene un sistema de ecuaciones lineales de orden  $N$  que al resolverlo se tienen los valores de  $C_i$ ,  $i = 1, \dots, N$ .

## Capítulo 3

# APLICACIÓN DE LOS MÉTODOS NUMÉRICOS A LA ECUACIÓN DE HELMHOLTZ EN 1-D

### 3.1. Método de Diferencias Finitas Implícito

Para obtener las ecuaciones en diferencia según el métodos de diferencias finitas se utilizan aproximaciones para  $p'$  en la ecuación (1.9)

$$p'(x_i) \approx \frac{p(x_i + h) - p(x_i - h)}{2h} \quad (3.1)$$

y para  $z'$  en (1.10)

$$z'(x_i) \approx \frac{z(x_i + h) - z(x_i - h)}{2h} \quad (3.2)$$

ambas aproximaciones de tres puntos centrada en  $x_i$  con  $i = 1, 2, \dots, n - 1$ . Denotando  $z(x_i)$  como  $z_i$  y de igual forma para la función  $p$  y luego reordenando se obtienen las ecuaciones en diferencia

$$2hz_i + p_{i-1} - p_{i+1} = 0 \quad (3.3)$$

$$-z_{i-1} + z_{i+1} + 2hk^2 p_i = 0 \quad (3.4)$$

Utilizando una aproximación para  $p'(x_0)$  con tres puntos hacia adelante

$$p'(x_0) \approx \frac{-3p(x_0) + 4p(x_1) - p(x_2)}{2h} \quad (3.5)$$

y tres puntos hacia atrás para  $z'(x_n)$

$$z'(x_n) \approx \frac{z(x_{n-2}) - 4z(x_{n-1}) + 3z(x_n)}{2h} \quad (3.6)$$

al tomar en cuenta las condiciones de frontera, la notación establecida anteriormente y que  $z_0 = p'_0$  tenemos

$$3p_0 - 4p_1 + p_2 = -2hz_0 = -2ihk \quad (3.7)$$

$$z_{n-2} - 4z_{n-1} + 3z_n + 2hk^2 p_n = z_{n-2} - 4z_{n-1} + (3 - 2ik)z_n = 0 \quad (3.8)$$

Para ilustrar el sistema resultante al aplicar éste método se presenta el caso en el que  $n = 4$ , obteniéndose

$$\begin{bmatrix} 3 & 0 & -4 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & a & 1 & 0 & 0 & 0 & 0 \\ 1 & b & 0 & 0 & -1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 & a & 1 & 0 & 0 \\ 0 & 0 & 1 & b & 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & -1 & 0 & 0 & a & 1 \\ 0 & 0 & 0 & 0 & 1 & b & 0 & \frac{-1}{c} \\ 0 & 0 & 0 & 1 & 0 & -4 & 0 & d \end{bmatrix} \begin{bmatrix} p_0 \\ z_1 \\ p_1 \\ z_2 \\ p_2 \\ z_3 \\ p_3 \\ z_4 \end{bmatrix} = \begin{bmatrix} -bc \\ c \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (3.9)$$

donde  $b = 2h$ ,  $a = bk^2$ ,  $c = ik$ , y  $d = 3 - bc$ . Se destaca el hecho que este es un sistema donde las variables  $p$  y  $z$  toman valores pertenecientes al campo de los números complejos.

## 3.2. Método de Elementos Finitos Galerkin Mixto

La formulación de elementos finitos Galerkin Mixto dada por (2.42) aplicada al sistema de ecuaciones diferenciales (1.9) y (1.10) con condiciones de frontera (1.11)

y (1.12), al considerar  $v = \begin{bmatrix} w \\ q \end{bmatrix}$ , quedaría de la siguiente forma

$$(Au, v) = \int_0^1 \begin{bmatrix} \bar{w} & \bar{q} \end{bmatrix} \begin{bmatrix} z - p' \\ z' + k^2 p \end{bmatrix} dx = 0 \quad (3.10)$$

que al desarrollar nos queda

$$\int_0^1 [\bar{w}(z - p') + \bar{q}(z' + k^2 p)] dx = 0 \quad (3.11)$$

Donde, al integrar el segundo término por partes, tenemos

$$\int_0^1 \bar{w}(z - p') dx + z(1)\bar{q}(1) - z(0)\bar{q}(0) - \int_0^1 z\bar{q}' dx + k^2 \int_0^1 p\bar{q} dx = 0 \quad (3.12)$$

Si utilizamos  $v = \varphi_i(x) \begin{bmatrix} 1 \\ 0 \end{bmatrix}$ ;  $i = 1, \dots, n$ , correspondería a utilizar en la ecuación (3.12) la sustitución  $w(x) = \varphi_i(x)$  y  $q(x) = 0$ , obtenemos

$$\int_0^1 \varphi_i(z - p') dx = 0; \quad i = 1, \dots, n \quad (3.13)$$

Es importante mencionar que  $i = 0$  no se considera en esta ecuación debido a que la condición de borde dada en la ecuación (1.11) conlleva a conocer el valor de  $z(0)$ , luego se retomará este punto.

Si utilizamos ahora  $v = \varphi_i(x) \begin{bmatrix} 0 \\ 1 \end{bmatrix}$ ;  $i = 0, \dots, n - 1$ , correspondería a utilizar en la ecuación (3.12) la sustitución  $w(x) = 0$  y  $q(x) = \varphi_i(x)$ , obtenemos la ecuación

$$z(1)\varphi_i(1) - z(0)\varphi_i(0) - \int_0^1 z\varphi_i' dx + k^2 \int_0^1 p\varphi_i dx = 0; \quad i = 0, \dots, n - 1 \quad (3.14)$$

Sustituyendo las condiciones de borde (1.11) y (1.12) en esta última ecuación

$$ikp(1)\varphi_i(1) - ik\varphi_i(0) - \int_0^1 z\varphi'_i dx + k^2 \int_0^1 p\varphi_i dx = 0; \quad i = 0, \dots, n-1 \quad (3.15)$$

En la ecuación (3.13) y posteriormente en la (3.15) se procederá a sustituir  $z$  y  $p$  según el método de elemento finito ya descrito, así

$$z(x) = \sum_{k=0}^n z_k \varphi_k(x) \quad (3.16)$$

$$p(x) = \sum_{k=0}^n p_k \varphi_k(x) \quad (3.17)$$

dando como resultado

$$\int_0^1 \varphi_i \left( \sum_{k=0}^n z_k \varphi_k - \sum_{k=0}^n p_k \varphi'_k \right) dx = 0; \quad i = 1, \dots, n \quad (3.18)$$

aplicando propiedades básicas de las integrales tenemos

$$\sum_{k=0}^n z_k \int_0^1 \varphi_k \varphi_i dx - \sum_{k=0}^n p_k \int_0^1 \varphi'_k \varphi_i dx = 0; \quad i = 1, \dots, n \quad (3.19)$$

Es necesario recordar que las constantes  $z_k$  y  $p_k$  son los valores de las funciones  $z$  y  $p$  en el nodo  $x_k$  (esto se verifica fácilmente observando el comportamiento de las ecuaciones (3.16) y (3.17) en el nodo  $x = x_i$ ), las cuales serán los valores que deseamos calcular para obtener la solución aproximada para ambas funciones.

Analicemos ahora la ecuación (3.19) para diferentes valores de  $i$

Si  $i = j$  con  $j = 1, \dots, n-1$

$$\begin{aligned} & z_{j-1} \int_0^1 \varphi_{j-1} \varphi_j dx + z_j \int_0^1 \varphi_j^2 dx + z_{j+1} \int_0^1 \varphi_{j+1} \varphi_j dx - \\ & p_{j-1} \int_0^1 \varphi'_{j-1} \varphi_j dx - p_j \int_0^1 \varphi'_j \varphi_j dx - p_{j+1} \int_0^1 \varphi'_{j+1} \varphi_j dx = 0 \end{aligned} \quad (3.20)$$

Sustituyendo los valores de las integrales según los cálculos realizados en el Apéndice G

$$\frac{h}{6} z_{j-1} + \frac{2h}{3} z_j + \frac{h}{6} z_{j+1} + \frac{1}{2} p_{j-1} - \frac{1}{2} p_{j+1} = 0 \quad (3.21)$$

Si  $i = n$

$$z_{n-1} \int_0^1 \varphi_{n-1} \varphi_n dx + z_n \int_0^1 \varphi_n^2 dx - p_{n-1} \int_0^1 \varphi'_{n-1} \varphi_n dx - p_n \int_0^1 \varphi'_n \varphi_n dx = 0 \quad (3.22)$$

Simplificando

$$\frac{h}{6} z_{n-1} + \frac{h}{3} z_n + \frac{1}{2} p_{n-1} - \frac{1}{2} p_n = 0 \quad (3.23)$$

Procedemos de igual forma con la ecuación (3.15), obteniendo

$$ik \sum_{k=0}^n p_k \varphi_k(1) \varphi_i(1) - ik \varphi_i(0) - \int_0^1 \left( \sum_{k=0}^n z_k \varphi_k \varphi'_i - k^2 \sum_{k=0}^n p_k \varphi_k \varphi_i \right) dx = 0 \quad (3.24)$$

aplicando algunas propiedades básicas de las integrales tenemos

$$ik \varphi_i(1) \sum_{k=0}^n p_k \varphi_k(1) - ik \varphi_i(0) - \sum_{k=0}^n z_k \int_0^1 \varphi_k \varphi'_i dx + k^2 \sum_{k=0}^n p_k \int_0^1 \varphi_k \varphi_i dx = 0 \quad (3.25)$$

Si  $i = 0$

$$-ik - z_0 \int_0^1 \varphi_0 \varphi'_0 dx - z_1 \int_0^1 \varphi_1 \varphi'_0 dx + k^2 p_0 \int_0^1 \varphi_0^2 dx + k^2 p_1 \int_0^1 \varphi_1 \varphi_0 dx = 0 \quad (3.26)$$

Sustituyendo los valores de las integrales, tenemos

$$-ik + \frac{1}{2} z_0 + \frac{1}{2} z_1 + k^2 \frac{h}{3} p_0 + k^2 \frac{h}{6} p_1 = 0 \quad (3.27)$$

Si  $i = j$  con  $j = 1, \dots, n-1$

$$\begin{aligned} & -z_{j-1} \int_0^1 \varphi_{j-1} \varphi'_j dx - z_j \int_0^1 \varphi_j \varphi'_j dx - z_{j+1} \int_0^1 \varphi_{j+1} \varphi'_j dx + \\ & k^2 p_{j-1} \int_0^1 \varphi_{j-1} \varphi_j dx + k^2 p_j \int_0^1 \varphi_j^2 dx + k^2 p_{j+1} \int_0^1 \varphi_{j+1} \varphi_j dx = 0 \end{aligned} \quad (3.28)$$

Simplificando

$$-\frac{1}{2} z_{j-1} + \frac{1}{2} z_{j+1} + k^2 \frac{h}{6} p_{j-1} + k^2 \frac{2h}{3} p_j + k^2 \frac{h}{6} p_{j+1} = 0 \quad (3.29)$$

Las ecuaciones (3.27), (3.21), (3.29) y (3.23) conforman un sistema de ecuaciones lineales de orden  $2n$ , este orden de colocar las ecuaciones es importante para

generar una matriz en banda, lo cual es deseado. Al conformar el sistema de ecuaciones se debe sustituir el valor de  $z_0$  según la condición de borde (1.11), de igual forma se sustituye  $p_n = \frac{z_n}{ik}$ . Para ilustrar el sistema resultante al aplicar este método se presenta el caso en el que  $n = 4$ , obteniéndose

$$\begin{bmatrix} 2c & \frac{1}{2} & c & 0 & 0 & 0 & 0 & 0 \\ \frac{1}{2} & 4a & 0 & a & -\frac{1}{2} & 0 & 0 & 0 \\ c & 0 & 4c & \frac{1}{2} & c & 0 & 0 & 0 \\ 0 & a & \frac{1}{2} & 4a & 0 & a & -\frac{1}{2} & 0 \\ 0 & -\frac{1}{2} & c & 0 & 4c & \frac{1}{2} & c & 0 \\ 0 & 0 & 0 & a & \frac{1}{2} & 4a & 0 & e \\ 0 & 0 & 0 & -\frac{1}{2} & c & 0 & 4c & g \\ 0 & 0 & 0 & 0 & 0 & a & \frac{1}{2} & l \end{bmatrix} \begin{bmatrix} p_0 \\ z_1 \\ p_1 \\ z_2 \\ p_2 \\ z_3 \\ p_3 \\ z_4 \end{bmatrix} = \begin{bmatrix} \frac{1}{2}f \\ -af \\ \frac{1}{2}f \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (3.30)$$

donde  $a = \frac{h}{6}$ ,  $b = k^2$ ,  $c = ab$ ,  $f = ik$ ,  $g = \frac{1}{2} + \frac{c}{f}$ ,  $e = a - \frac{1}{2f}$  y  $l = 2a - \frac{1}{2f}$ .

### 3.3. Método de Elementos Finitos Mínimos Cuadrados para Sistemas de Primer Orden

Dada la ecuación diferencial lineal  $Au = f$ , sujeta a las condiciones de frontera  $Bu = g$ . Consideraremos el problema homogéneo, es decir,  $f = 0$  y  $g = 0$ . Se desea minimizar el funcional  $I(u) = \frac{1}{2}\|Au\|_0^2 + \frac{1}{2}\|Bu\|_\gamma^2$  para ello se sabe que una condición necesaria es

$$\lim_{t \rightarrow 0} \frac{dI(u + tv)}{dt} = 0 \quad (3.31)$$

Tenemos que

$$I(u + tv) = \frac{1}{2}(A(u + tv), A(u + tv)) + \frac{1}{2}\langle B(u + tv), B(u + tv) \rangle_\gamma \quad (3.32)$$

como ya se ha dicho  $A$  y  $B$  son operadores diferenciales lineales, aplicando las propiedades del producto interno y luego derivando respecto a  $t$ , se obtiene

$$\begin{aligned} \frac{dI(u + tv)}{dt} = & \frac{1}{2}[(Au, Av) + (Av, Au) + 2t(Av, Av)] + \\ & \frac{1}{2}[\langle Bu, Bv \rangle_\gamma + \langle Bv, Bu \rangle_\gamma + 2t\langle Bv, Bv \rangle_\gamma] \end{aligned} \quad (3.33)$$

luego al aplicar el límite cuando  $t$  tiende a cero e igualando a cero

$$\lim_{t \rightarrow 0} \frac{dI(u + tv)}{dt} = \frac{1}{2}[(Au, Av) + (Av, Au)] + \frac{1}{2}[\langle Bu, Bv \rangle_\gamma + \langle Bv, Bu \rangle_\gamma] = 0 \quad (3.34)$$

Basándonos en uno de los axiomas para el producto interno de un espacio vectorial complejo, el cual establece que  $(x, y) = \overline{(y, x)}$  y como  $(x, y) + \overline{(x, y)} = 2\text{Re}(x, y)$  se concluye que

$$\text{Re}(Au, Av) + \text{Re}\langle Bu, Bv \rangle_\gamma = 0 \quad (3.35)$$

Si se utiliza  $v = \varphi_i(x) \begin{bmatrix} 1 \\ 0 \end{bmatrix}$ , en el primer término de la ecuación (3.35) aparecerán de forma explícita  $z^r$  y  $q^r$  y en el segundo término  $z^{im}$  y  $q^{im}$ , donde los superíndices  $r$  y  $im$  hacen referencia a la parte real e imaginaria de las variables  $z$  y  $q$  según corresponda. Por ello, se hace necesario utilizar una base para las funciones  $v$  que incluyan parte real y parte imaginaria, esto no hacía falta en el método Galerkin Mixto. En los próximos párrafos, se utilizarán como elementos de una base de funciones imaginarias, donde se encuentre  $v$ , a las siguientes funciones:

$$v = \varphi_m(x) \begin{bmatrix} 1 \\ 0 \end{bmatrix}, v = \varphi_m(x) \begin{bmatrix} 0 \\ 1 \end{bmatrix}, v = \varphi_m(x) \begin{bmatrix} i \\ 0 \end{bmatrix} \text{ y } v = \varphi_m(x) \begin{bmatrix} 0 \\ i \end{bmatrix}; m = 1, \dots, n \text{ e } i = \sqrt{-1}$$

Nos dedicaremos inicialmente a desarrollar el primer término de la ecuación (3.35) para el caso que nos compete. Recordemos que se deseaba resolver el sistema de ecuaciones diferenciales dado por (1.9)-(1.12), así al sustituir a  $A$ ,  $u$  y  $v$  de

la misma forma como se hizo en la formulación del método de elemento finito Galerkin Mixto en la ecuación

$$Re(Au, Av) = Re\left(\int_0^1 (\overline{Av})^t Audx\right) \quad (3.36)$$

y considerando  $v = \begin{bmatrix} w \\ q \end{bmatrix}$ , se obtiene

$$Re(Au, Av) = Re \int_0^1 \left( \overline{\begin{bmatrix} 1 & 0 \\ 0 & k^2 \end{bmatrix} \begin{bmatrix} w \\ q \end{bmatrix} + \begin{bmatrix} 0 & -1 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} w' \\ q' \end{bmatrix}} \right)^t \left( \begin{bmatrix} 1 & 0 \\ 0 & k^2 \end{bmatrix} \begin{bmatrix} z \\ p \end{bmatrix} + \begin{bmatrix} 0 & -1 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} z' \\ p' \end{bmatrix} \right) dx \quad (3.37)$$

que al simplificar nos queda

$$Re(Au, Av) = Re \int_0^1 [(\overline{w} - \overline{q'})(z - p') + (k^2 \overline{q} + \overline{w'})(z' + k^2 p)] dx \quad (3.38)$$

Si utilizamos  $v = \varphi_m(x) \begin{bmatrix} 1 \\ 0 \end{bmatrix}$ ;  $m = 1, \dots, n$ , o sea real puro, correspondería a utilizar en la ecuación (3.38) la sustitución  $w(x) = \varphi_m(x)$  y  $q(x) = 0$ , obtenemos

$$Re(Au, Av) = Re \int_0^1 [\varphi_m(x)(z - p') + \varphi'_m(x)(z' + k^2 p)] dx \quad (3.39)$$

Si ahora usamos  $v = \varphi_m(x) \begin{bmatrix} 0 \\ 1 \end{bmatrix}$ ;  $m = 1, \dots, n$ , otra vez real puro, correspondería a utilizar en la ecuación (3.38) la sustitución  $q(x) = \varphi_m(x)$  y  $w(x) = 0$ , se obtiene

$$Re(Au, Av) = Re \int_0^1 [-\varphi'_m(x)(z - p') + k^2 \varphi_m(x)(z' + k^2 p)] dx \quad (3.40)$$

Si  $v = \varphi_m(x) \begin{bmatrix} i \\ 0 \end{bmatrix}$ ;  $m = 1, \dots, n$ , o sea imaginario puro, correspondería a utilizar en la ecuación (3.38) la sustitución  $q(x) = i\varphi_m(x)$  y  $w(x) = 0$ , obteniéndose

$$Re(Au, Av) = Re(-i \int_0^1 [\varphi_m(x)(z - p') + \varphi'_m(x)(z' + k^2 p)] dx) \quad (3.41)$$

Si  $v = \varphi_m(x) \begin{bmatrix} 0 \\ i \end{bmatrix}$ ;  $m = 1, \dots, n$ , otra vez imaginario puro, correspondería a utilizar en la ecuación (3.38) la sustitución  $q(x) = i\varphi_m(x)$  y  $w(x) = 0$ , obteniéndose

$$Re(Au, Av) = Re(-i \int_0^1 [-\varphi'_m(x)(z - p') + k^2 \varphi_m(x)(z' + k^2 p)] dx) \quad (3.42)$$

Obsérvese que las ecuaciones (3.39) y (3.41) difieren solamente en un factor a saber  $-i$ , dando como resultado ecuaciones similares, las cuales se diferencian sólo en el superíndice que serían  $r$  y  $im$  respectivamente. Lo mismo ocurre en las ecuaciones (3.40) y (3.42). Por ello sólo se desarrollarán las ecuaciones (3.39) y (3.40)

Comencemos con la ecuación (3.39). Primero se sustituye  $z$  y  $p$  según las ecuaciones (3.16) y (3.17)

$$Re(Au, Av) = Re \int_0^1 [\varphi_m(\sum_{k=0}^n z_k \varphi_k - \sum_{k=0}^n p_k \varphi'_k) + \varphi'_m(\sum_{k=0}^n z_k \varphi'_k + k^2 \sum_{k=0}^n p_k \varphi_k)] dx \quad (3.43)$$

Aplicando propiedades de las integrales

$$Re(Au, Av) = Re[\sum_{k=0}^n z_k (\int_0^1 \varphi_m \varphi_k dx + \int_0^1 \varphi'_m \varphi'_k dx) - \sum_{k=0}^n p_k (\int_0^1 \varphi_m \varphi'_k dx - k^2 \int_0^1 \varphi'_m \varphi_k dx)] \quad (3.44)$$

Ahora se analizan los diferentes valores de  $m$

Si  $m = j$ ;  $j = 1, \dots, n-1$

$$\begin{aligned} Re(Au, Av) = & Re[z_{j-1} \int_0^1 (\varphi_{j-1} \varphi_j + \varphi'_{j-1} \varphi'_j) dx] + Re[z_j \int_0^1 (\varphi_j^2 + (\varphi'_{j-1})^2) dx] \\ & + Re[z_{j+1} \int_0^1 (\varphi_{j+1} \varphi_j + \varphi'_{j+1} \varphi'_j) dx] + Re[p_{j-1} \int_0^1 (k^2 \varphi_{j-1} \varphi'_j - \varphi'_{j-1} \varphi_j) dx] \\ & + Re[p_j \int_0^1 (k^2 - 1) \varphi_j \varphi'_j dx] + Re[p_{j+1} \int_0^1 (k^2 \varphi_{j+1} \varphi'_j - \varphi'_{j+1} \varphi_j) dx] \end{aligned} \quad (3.45)$$

Evaluable las integrales, las cuales son todas reales y siendo el superíndice  $r$  el identificador para la parte real de la variable a la cual está aplicada

$$Re(Au, Av) = (\frac{h}{6} - \frac{1}{h})z_{j-1}^r + (\frac{2h}{3} + \frac{2}{h})z_j^r + (\frac{h}{6} - \frac{1}{h})z_{j+1}^r + \frac{(1+k^2)}{2}p_{j-1}^r - \frac{(1+k^2)}{2}p_{j+1}^r \quad (3.46)$$

Si  $m = n$

$$Re(Au, Av) = Re[z_{n-1} \int_0^1 (\varphi_{n-1}\varphi_n + \varphi'_{n-1}\varphi'_n)dx] + Re[z_n \int_0^1 (\varphi_n^2 + (\varphi'_n)^2)dx] \\ + Re[p_{n-1} \int_0^1 (k^2\varphi_{n-1}\varphi'_n - \varphi'_{n-1}\varphi_n)dx] + Re[p_n \int_0^1 (k^2 - 1)\varphi_n\varphi'_n dx] \quad (3.47)$$

Evaluable las integrales, nos queda

$$Re(Au, Av) = (\frac{h}{6} - \frac{1}{h})z_{n-1}^r + (\frac{h}{3} + \frac{1}{h})z_n^r + \frac{(1+k^2)}{2}p_{n-1}^r + \frac{(k^2-1)}{2}p_n^r \quad (3.48)$$

Continuaremos ahora con la ecuación (3.40), sustituyéndose las ecuaciones (3.16) y (3.17)

$$Re(Au, Av) = -Re \int_0^1 [\varphi'_m (\sum_{k=0}^n z_k \varphi_k - \sum_{k=0}^n p_k \varphi'_k) + \varphi_m (k^2 \sum_{k=0}^n z_k \varphi'_k + k^4 \sum_{k=0}^n p_k \varphi_k)] dx \quad (3.49)$$

Aplicando propiedades de las integrales

$$Re(Au, Av) = Re[\sum_{k=0}^n z_k (k^2 \int_0^1 \varphi_m \varphi'_k dx - \int_0^1 \varphi'_m \varphi_k dx) \\ + \sum_{k=0}^n p_k (\int_0^1 \varphi'_m \varphi'_k dx + k^4 \int_0^1 \varphi_m \varphi_k dx)] \quad (3.50)$$

Ahora se analizan los diferentes valores de  $m$

Si  $m = 0$

$$Re(Au, Av) = Re[z_0 \int_0^1 (k^2 - 1)\varphi'_0 \varphi_0 dx] + Re[z_1 \int_0^1 (k^2 \varphi'_1 \varphi_0 - \varphi_1 \varphi'_0) dx] \\ + Re[p_0 \int_0^1 (k^4 \varphi_0^2 + (\varphi'_0)^2) dx] + Re[p_1 \int_0^1 (k^4 \varphi_1 \varphi_0 + \varphi'_1 \varphi'_0) dx] \quad (3.51)$$

Evalutando las integrales

$$Re(Au, Av) = \frac{(1-k^2)}{2}z_0^r + \frac{(1+k^2)}{2}z_1^r + (\frac{k^4h}{3} + \frac{1}{h})p_0^r + (\frac{k^4h}{6} - \frac{1}{h})p_1^r \quad (3.52)$$

Si  $m = j$ ;  $j = 1, \dots, n-1$

$$\begin{aligned} Re(Au, Av) = & Re[z_{j-1} \int_0^1 (k^2 \varphi'_{j-1} \varphi_j - \varphi_{j-1} \varphi'_j) dx] + Re[z_j \int_0^1 (k^2 - 1) \varphi'_j \varphi_j dx] \\ & + Re[z_{j+1} \int_0^1 (k^2 \varphi'_{j+1} \varphi_j - \varphi_{j+1} \varphi'_j) dx] + Re[p_{j-1} \int_0^1 (k^4 \varphi_{j-1} \varphi_j(x) + \varphi'_{j-1} \varphi'_j) dx] \\ & + Re[p_j \int_0^1 (k^4 \varphi_j^2 + (\varphi'_j)^2) dx] + Re[p_{j+1} \int_0^1 (k^4 \varphi_{j+1} \varphi_j + \varphi'_{j+1} \varphi'_j) dx] \end{aligned} \quad (3.53)$$

Evalutando las integrales

$$\begin{aligned} Re(Au, Av) = & -\frac{(1+k^2)}{2}z_{j-1}^r + \frac{(1+k^2)}{2}z_{j+1}^r + (\frac{k^4h}{6} - \frac{1}{h})p_{j-1}^r + (\frac{2k^4h}{3} + \frac{2}{h})p_j^r + (\frac{k^4h}{6} - \frac{1}{h})p_{j+1}^r \end{aligned} \quad (3.54)$$

Si  $m = n$

$$\begin{aligned} Re(Au, Av) = & Re[z_{n-1} \int_0^1 (k^2 \varphi'_{n-1} \varphi_n - \varphi_{n-1} \varphi'_n) dx] + Re[z_n \int_0^1 (k^2 - 1) \varphi_n \varphi'_n dx] \\ & + Re[p_{n-1} \int_0^1 (k^4 \varphi_{n-1} \varphi_n + \varphi'_{n-1} \varphi'_n) dx] + Re[p_n \int_0^1 ((\varphi'_n)^2 + k^4 (\varphi_n)^2) dx] \end{aligned} \quad (3.55)$$

Evalutando las integrales, nos queda

$$Re(Au, Av) = -\frac{(1+k^2)}{2}z_{n-1}^r + \frac{(k^2-1)}{2}z_n^r + (\frac{k^4h}{6} - \frac{1}{h})p_{n-1}^r + (\frac{k^4h}{3} + \frac{1}{h})p_n^r \quad (3.56)$$

Para simplificar el segundo término de la ecuación (3.35) no se considerará la condición de borde dada por (1.11) ya que cuando generemos el sistema de ecuaciones lineales se va a sustituir el valor de la incógnita que representa a  $p'(0)$  por el valor  $ik$ .

Si en la ecuación

$$Re\langle Bu, Bv \rangle = Re[(\overline{Bv})^t Bu] \quad (3.57)$$

sustituimos  $v = \psi_n(1) \begin{bmatrix} 1 \\ 0 \end{bmatrix}$  se obtiene

$$Re\langle Bu, Bv \rangle = Re\left[\overline{\left([1 \quad -ik] \begin{bmatrix} \psi_n(1) \\ 0 \end{bmatrix}\right)^t} [1 \quad -ik] \begin{bmatrix} z_n \\ p_n \end{bmatrix}\right] \quad (3.58)$$

Al sustituir  $\psi_n(1) = 1$  en la ecuación anterior y simplificar, obtenemos

$$Re\langle Bu, Bv \rangle = Re[z_n - ikp_n] = z_n^r + kp_n^{im} \quad (3.59)$$

Si  $v = \psi_n(1) \begin{bmatrix} 0 \\ 1 \end{bmatrix}$

$$Re\langle Bu, Bv \rangle = Re\left[\overline{\left([1 \quad -ik] \begin{bmatrix} 0 \\ \psi_n(1) \end{bmatrix}\right)^t} [1 \quad -ik] \begin{bmatrix} z_n \\ p_n \end{bmatrix}\right] \quad (3.60)$$

Simplificando

$$Re\langle Bu, Bv \rangle = Re[ik(z_n - ikp_n)] = -kz_n^{im} + k^2p_n^r \quad (3.61)$$

Si  $v = \psi_n(1) \begin{bmatrix} i \\ 0 \end{bmatrix}$

$$Re\langle Bu, Bv \rangle = Re\left[\overline{\left([1 \quad -ik] \begin{bmatrix} i\psi_n(1) \\ 0 \end{bmatrix}\right)^t} [1 \quad -ik] \begin{bmatrix} z_n \\ p_n \end{bmatrix}\right] \quad (3.62)$$

Simplificando

$$Re\langle Bu, Bv \rangle = Re[-iz_n - kp_n] = z_n^{im} - kp_n^r \quad (3.63)$$

$$\text{Si } v = \psi_n(1) \begin{bmatrix} 0 \\ i \end{bmatrix}$$

$$Re\langle Bu, Bv \rangle = Re\left[\overline{\begin{bmatrix} 1 & -ik \end{bmatrix} \begin{bmatrix} 0 \\ i\psi_n(1) \end{bmatrix}}\right]^t \begin{bmatrix} 1 & -ik \end{bmatrix} \begin{bmatrix} z_n \\ p_n \end{bmatrix} \quad (3.64)$$

Simplificando

$$Re\langle Bu, Bv \rangle = Re[k(z_n - ikp_n)] = kz_n^r + k^2 p_n^{im} \quad (3.65)$$

Para conformar el sistema de ecuaciones lineales debemos sustituir los dos términos de la ecuación (3.35), donde el segundo término es cero, excepto en la frontera para  $x = 1$ . Recuerde que debe utilizarse tanto las sustituciones de  $v$  real e imaginaria. Al hacer dichas sustituciones se obtiene el sistema de ecuaciones lineales  $Ku = b_1$ , donde  $K = \begin{bmatrix} C & D \\ D^t & C \end{bmatrix}$ ,  $u = \begin{bmatrix} u^r \\ u^{im} \end{bmatrix}$  y  $b_1$  sólo tiene los tres primeros elementos no nulos, los cuales se obtienen de evaluar  $z_0^{im} = ik$  y pasarlo al miembro de la derecha de la ecuación como término independiente en las tres primeras ecuaciones.

A continuación para ilustrar un poco mejor como están conformados cada uno de los elementos de dicho sistema de ecuaciones lineales se presenta el caso para  $n = 4$ . En este caso

$$C = \begin{bmatrix} e & d & f & 0 & 0 & 0 & 0 & 0 & 0 \\ d & 2a & 0 & b & -d & 0 & 0 & 0 & 0 \\ f & 0 & 2e & d & f & 0 & 0 & 0 & 0 \\ 0 & b & d & 2a & 0 & b & -d & 0 & 0 \\ 0 & -d & f & 0 & 2e & d & f & 0 & 0 \\ 0 & 0 & 0 & b & d & 2a & 0 & b & -d \\ 0 & 0 & 0 & -d & f & 0 & 2e & d & f \\ 0 & 0 & 0 & 0 & 0 & b & d & a+1 & -c \\ 0 & 0 & 0 & 0 & 0 & -d & f & -c & e+k^2 \end{bmatrix} \quad (3.66)$$

donde  $a = \frac{h}{3} + \frac{1}{h}$ ,  $b = \frac{h}{6} - \frac{1}{h}$ ,  $c = \frac{1-k^2}{2}$ ,  $d = \frac{1+k^2}{2}$ ,  $e = \frac{k^4 h}{3} + \frac{1}{h}$  y  $f = \frac{k^4 h}{6} - \frac{1}{h}$ .

Se puede observar que  $C$  es una matriz simétrica de orden  $(2n+1)$  y los todos los elementos de ella son obtenidos sólo a partir de  $Re(Au, Av)$  menos los dos últimos elementos de la diagonal principal que también están formados por los términos 1 y  $k^2$  que provienen de  $Re\langle Bu, Bv \rangle_\gamma$ . la matriz  $D$  sólo tiene dos elementos no nulos, los cuales corresponden a las posiciones  $(2n, 2n+1)$  y  $(2n+1, 2n)$ , siendo los mismos  $k$  y  $-k$

$$u^r = \begin{bmatrix} p_0^r \\ z_1^r \\ p_1^r \\ z_2^r \\ p_2^r \\ z_3^r \\ p_3^r \\ z_4^r \\ p_4^r \end{bmatrix}, \quad u^{im} = \begin{bmatrix} p_0^{im} \\ z_1^{im} \\ p_1^{im} \\ z_2^{im} \\ p_2^{im} \\ z_3^{im} \\ p_3^{im} \\ z_4^{im} \\ p_4^{im} \end{bmatrix} \quad y \quad b_1 = \begin{bmatrix} -ick \\ -ibk \\ idk \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \quad (3.67)$$

La utilización alternada de las variables  $z$  y  $p$  es necesaria para que la matriz  $A$  sea simétrica y en banda.

# Capítulo 4

## Resultados y Conclusiones

### 4.1. Análisis del Orden de Convergencia

Para hacer un análisis del orden de convergencia de los métodos desarrollados en este trabajo, los mismos se aplican tomando en consideración una discretización uniforme del intervalo  $[0, 1]$  e ir incrementando la cantidad de subintervalos de forma sucesiva, para ello en los programas hechos en *MATLAB* se define una variable a la que se denominó *factor*, así el número inicial de intervalos  $n$  es *factor* y se va incrementando de *factor* en *factor* hasta que alcanza el valor de *orden\_matriz\_graficas* al cuadrado por *factor*. Luego se calculan los errores cometidos en la aproximación para cada valor de  $n$ ; pudiéndose, a través de una regresión lineal, calcular la pendiente de dicha recta que representa el orden de convergencia de cada uno de los métodos, esto es ya que el error está acotado, o sea  $\|Error\| \leq Ch^r$ , donde  $r$  es el denominado orden de convergencia. En este trabajo la medida del error se calcula basándose en la norma dos ( $\|\cdot\|_2$ ) de la diferencia entre la solución analítica y la solución numérica.

En todos los métodos se presenta una figura que contiene nueve gráficas de convergencia que corresponden a los valores de  $k = 1$  a  $k = 9$ . También se presentan cuadros para diversos valores de *factor* donde se muestran tanto el orden de convergencia para la variable primaria  $p$  como para la variable secundaria  $z$ , para diversos valores de  $k$ . Además se presentan en la última columna de dichos cuadros el valor del tiempo, en segundos, empleado para realizar los cálculos.

#### 4.1.1. Método de Diferencias Finitas Implícito

En las figuras 4.1, 4.2, 4.3 y 4.4 se presentan las curvas que representan a  $-\log(error)$  en función de  $-\log(h)$  para las variable  $z$   $p$  que son rectas casi paralelas teniendo pendientes aproximadamente igual a dos, veánse los cuadros 4.1, 4.2, 4.3 y 4.4, por lo que se concluye que el orden de convergencia para el error es cuadrático tanto para la variable primal  $p$  como para la secundaria  $z$ . En el anexo A se presenta el código en *MATLAB* generado para realizar dichos cálculos.

En base a los cuadros citados anteriormente se puede tambien concluir que a medida que se hacen discretizaciones más finas para un mismo valor de  $k$  no se observa una variación significativa del orden del error calculado.

#### 4.1.2. Método de Elementos Finitos Galerkin Mixto

En este caso se puede observar en las figuras 4.5, 4.6, 4.7 y 4.8 que las rectas que representan a  $-\log(error)$  en función de  $-\log(h)$  no son paralelas teniendo una pendiente aproximadamente igual a dos para la variable primal  $p$  e igual a uno

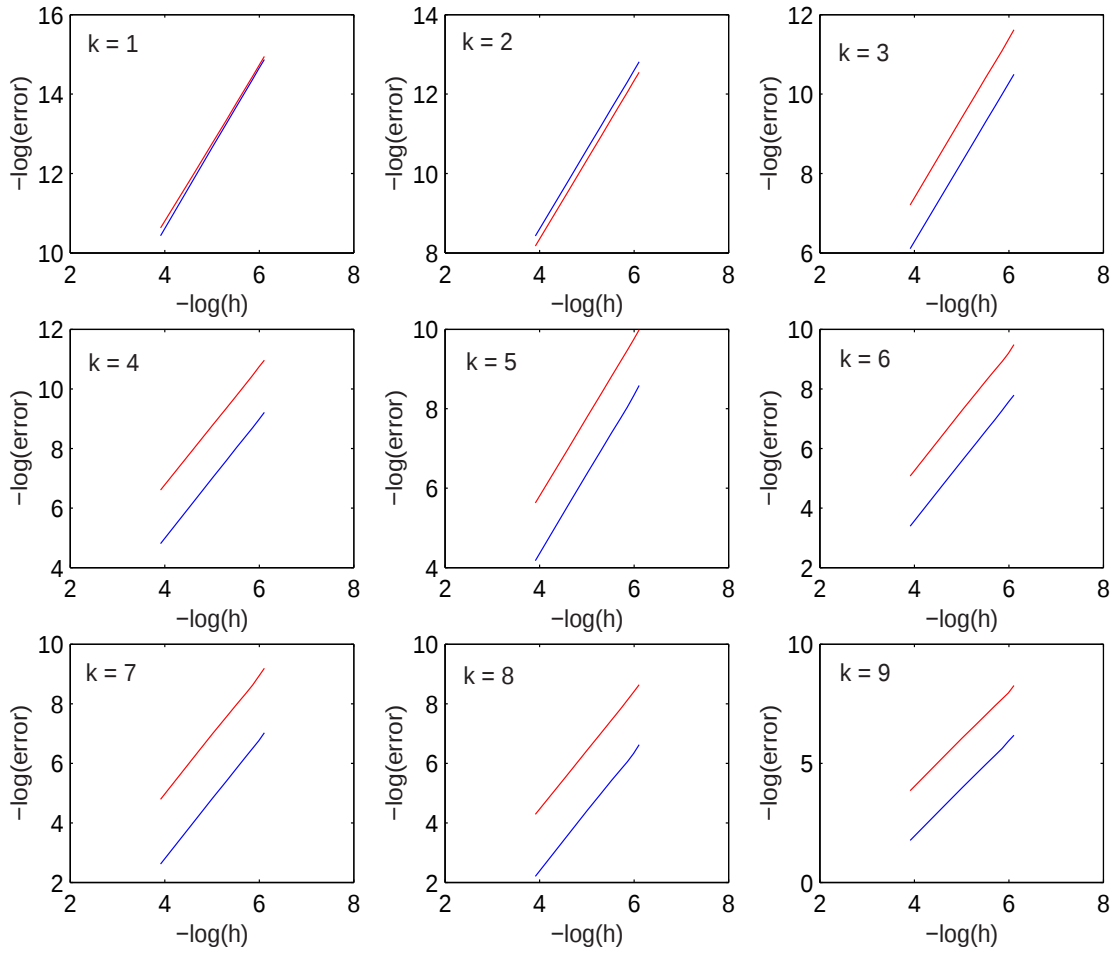


Figura 4.1: Error  $factor = 50$  (FDM)

| $k$ | $Orden(p)$ | $Orden(z)$ | $tiempo(s)$ |
|-----|------------|------------|-------------|
| 1   | 1.9667     | 2.0189     | 0.3440      |
| 2   | 1.9935     | 1.9953     | 0.3440      |
| 3   | 2.0008     | 1.9991     | 0.3280      |
| 4   | 1.9810     | 2.0012     | 0.3280      |
| 5   | 1.9800     | 1.9952     | 0.3280      |
| 6   | 1.9853     | 1.9995     | 0.3280      |
| 7   | 1.9806     | 1.9966     | 0.3280      |
| 8   | 1.9755     | 1.9880     | 0.3280      |
| 9   | 1.9822     | 1.9951     | 0.3280      |

Cuadro 4.1: Orden de Convergencia  $factor = 50$  (FDM)

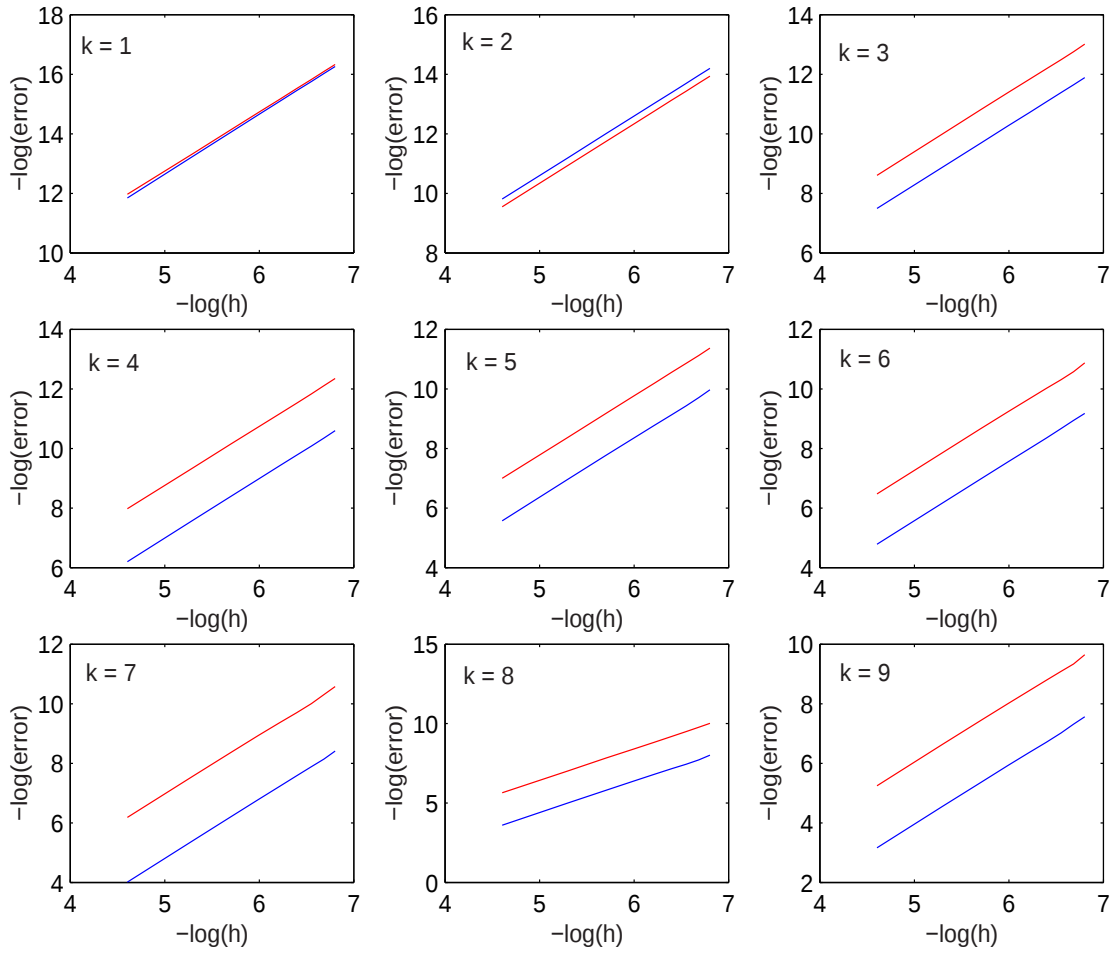


Figura 4.2: Error  $factor = 100$  (FDM)

| $k$ | $Orden(p)$ | $Orden(z)$ | $tiempo(s)$ |
|-----|------------|------------|-------------|
| 1   | 1.9835     | 2.0095     | 0.6720      |
| 2   | 1.9964     | 1.9971     | 0.6570      |
| 3   | 1.9969     | 1.9992     | 0.6560      |
| 4   | 1.9867     | 1.9978     | 0.6560      |
| 5   | 1.9873     | 1.9909     | 0.6560      |
| 6   | 1.9818     | 1.9975     | 0.6560      |
| 7   | 1.9780     | 1.9924     | 0.6720      |
| 8   | 1.9830     | 1.9816     | 0.6880      |
| 9   | 1.9786     | 1.9900     | 0.6870      |

Cuadro 4.2: Orden de Convergencia  $factor = 100$  (FDM)

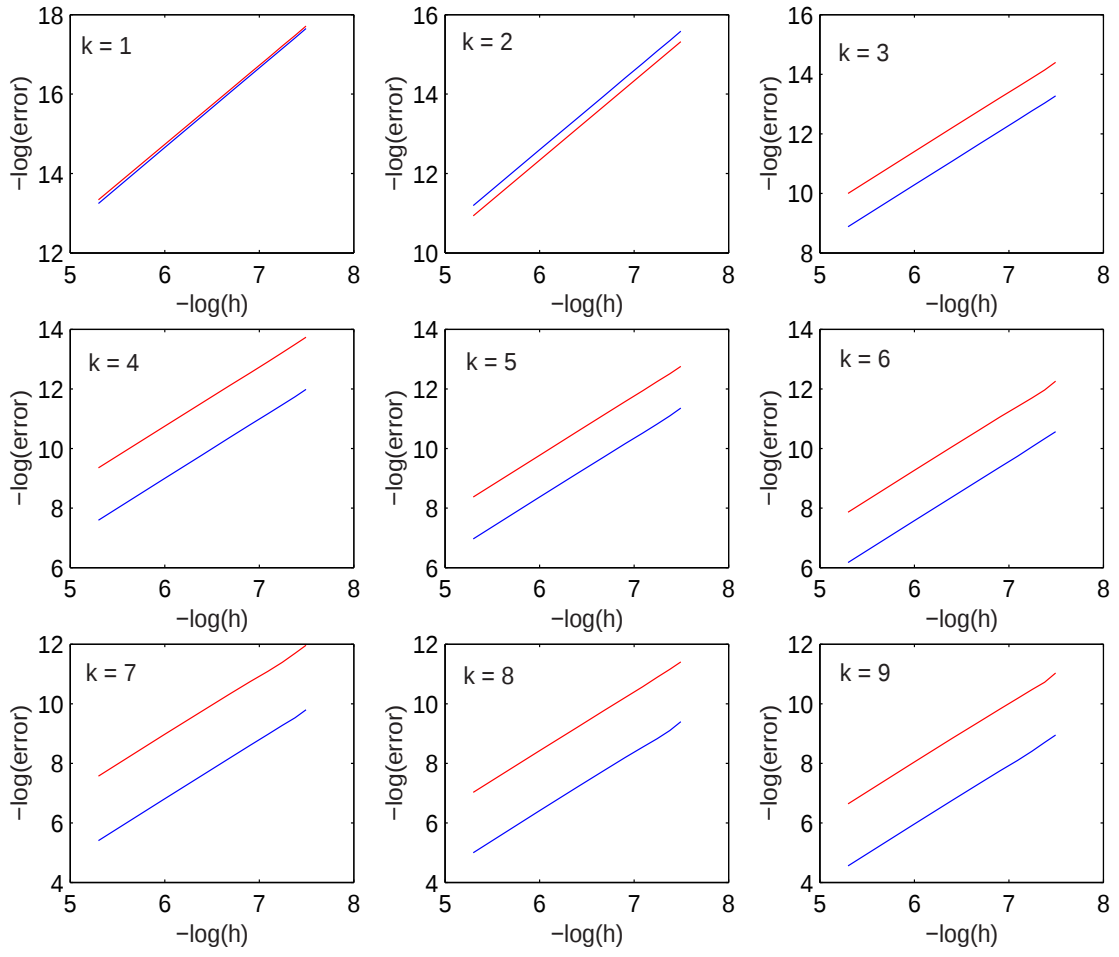


Figura 4.3: Error  $factor = 200$  (FDM)

| $k$ | $Orden(p)$ | $Orden(z)$ | $tiempo(s)$ |
|-----|------------|------------|-------------|
| 1   | 1.9918     | 2.0048     | 1.3750      |
| 2   | 1.9979     | 1.9982     | 1.3280      |
| 3   | 1.9950     | 1.9994     | 1.3590      |
| 4   | 1.9909     | 1.9963     | 1.3750      |
| 5   | 1.9912     | 1.9891     | 1.3280      |
| 6   | 1.9803     | 1.9972     | 1.3280      |
| 7   | 1.9785     | 1.9910     | 1.3430      |
| 8   | 1.9877     | 1.9790     | 1.3590      |
| 9   | 1.9774     | 1.9888     | 1.3600      |

Cuadro 4.3: Orden de Convergencia  $factor = 200$  (FDM)

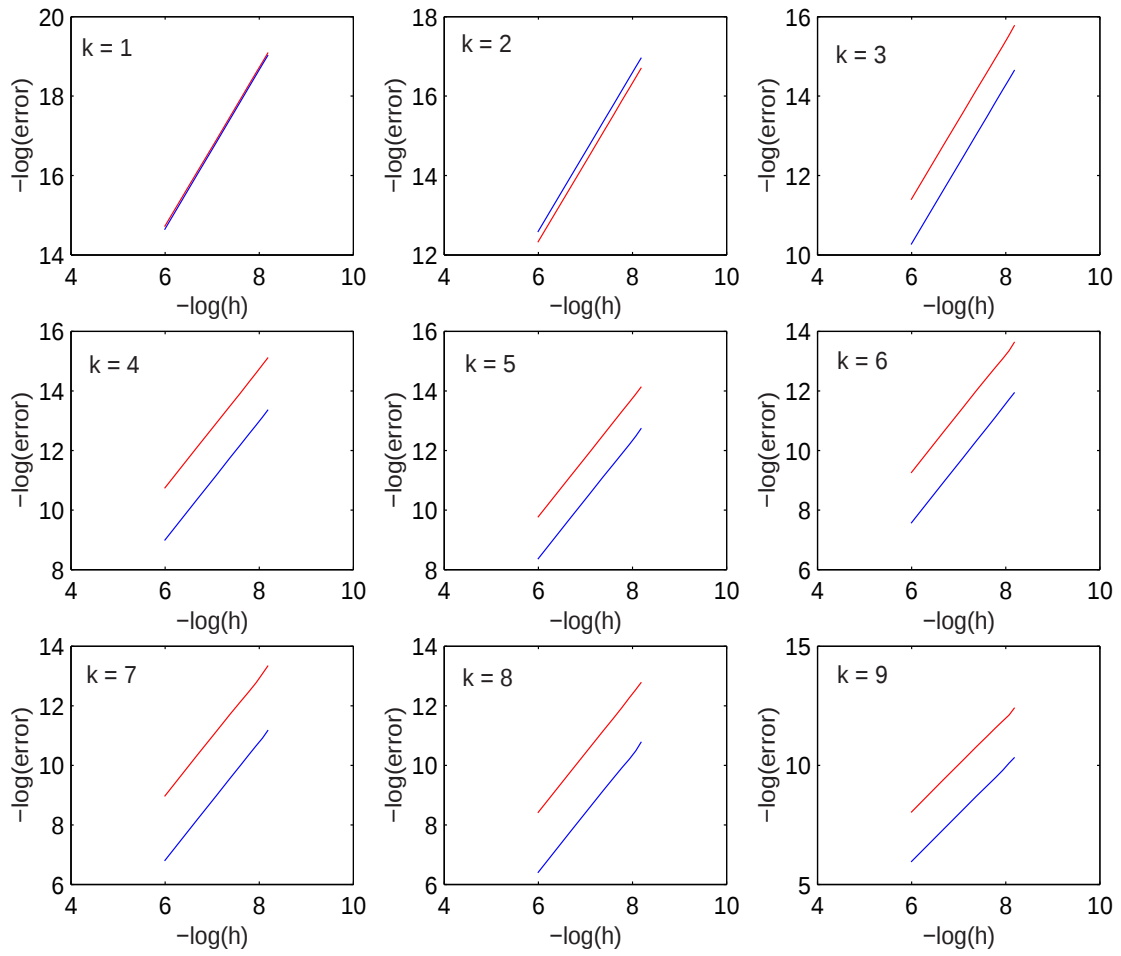


Figura 4.4: Error  $factor = 400$  (FDM)

| $k$ | $Orden(p)$ | $Orden(z)$ | $tiempo(s)$ |
|-----|------------|------------|-------------|
| 1   | 1.9959     | 2.0024     | 2.9540      |
| 2   | 1.9987     | 1.9989     | 2.8440      |
| 3   | 1.9940     | 1.9996     | 2.8130      |
| 4   | 1.9933     | 1.9956     | 2.8280      |
| 5   | 1.9932     | 1.9882     | 2.8590      |
| 6   | 1.9796     | 1.9972     | 2.9060      |
| 7   | 1.9792     | 1.9906     | 2.9220      |
| 8   | 1.9903     | 1.9778     | 2.8590      |
| 9   | 1.9769     | 1.9886     | 2.8750      |

Cuadro 4.4: Orden de Convergencia  $factor = 400$  (FDM)

para la variable secundaria  $z$ , comportándose este método de orden dos y orden uno respectivamente para cada variable. A medida que se hacen discretizaciones más finas para un mismo valor de  $k$  no se observan cambios significativos en el orden del error. Mención particular debe hacerse para el valor de  $k$  igual a ocho, el cual presenta un comportamiento irregular para el orden de convergencia para la variable primaria  $p$ , como se puede apreciar en la curvatura que se presenta en el lado derecho de cada una de sus gráficas en las figuras 4.5, 4.6, 4.7 y 4.8. En este caso, se observa un valor para el orden de convergencia que no sigue el comportamiento para los otros valores de  $k$ , para ello analícese los cuadros 4.5, 4.6, 4.7 y 4.8. Hasta el momento no sabemos justificar este hecho. En el anexo B se presenta el código en *MATLAB* generado para realizar dichos cálculos.

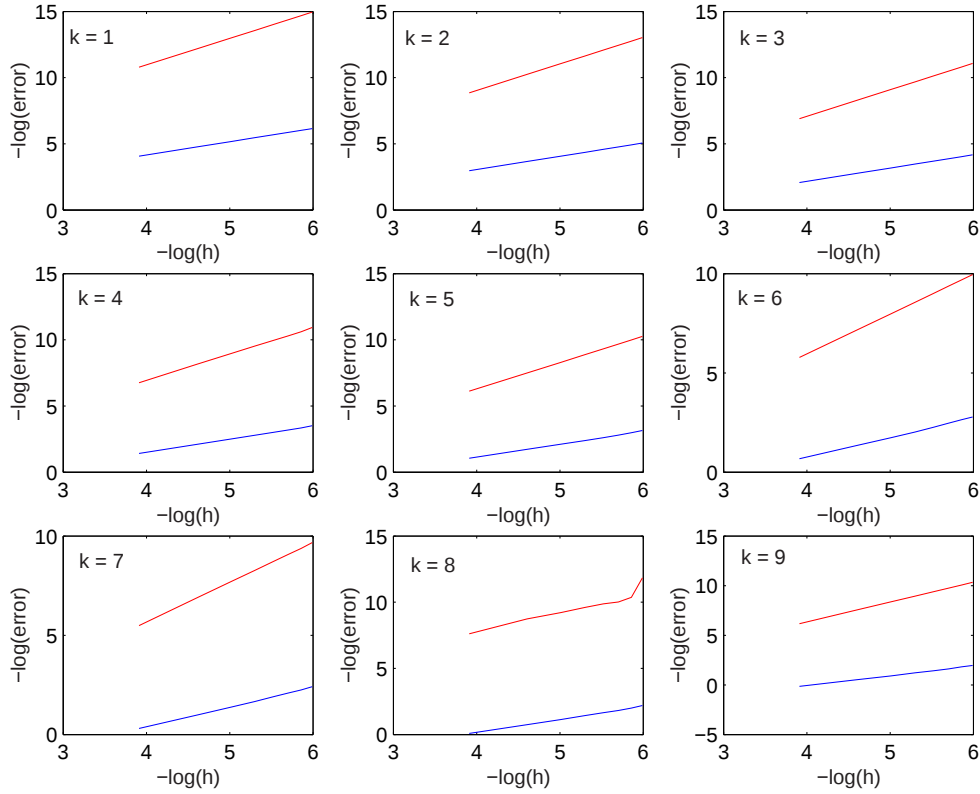


Figura 4.5: Error  $factor = 50$  (FEMGM)

| $k$ | $Orden(p)$ | $Orden(z)$ | $tiempo(s)$ |
|-----|------------|------------|-------------|
| 1   | 2.0014     | 0.9999     | 0.7650      |
| 2   | 2.0068     | 1.0034     | 0.6570      |
| 3   | 2.0054     | 1.0044     | 0.6560      |
| 4   | 1.9867     | 0.9963     | 0.6570      |
| 5   | 1.9931     | 0.9965     | 0.6560      |
| 6   | 2.0064     | 1.0121     | 0.6570      |
| 7   | 2.0005     | 1.0080     | 0.6560      |
| 8   | 1.6481     | 1.0001     | 0.6570      |
| 9   | 2.0036     | 1.0099     | 0.6560      |

Cuadro 4.5: Orden de Convergencia  $factor = 50$  (FEMGM)

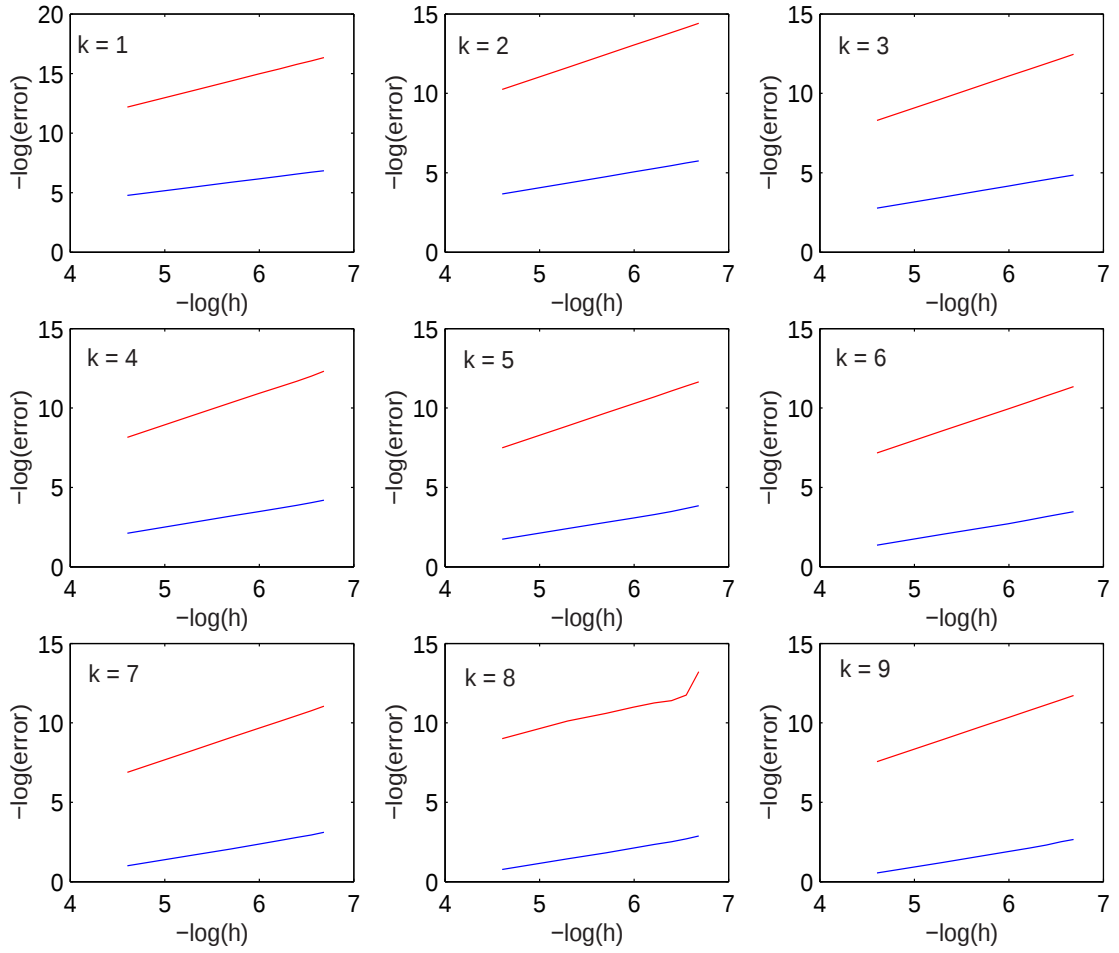


Figura 4.6: Error  $factor = 100$  (FEMGM)

| $k$ | $Orden(p)$ | $Orden(z)$ | $tiempo(s)$ |
|-----|------------|------------|-------------|
| 1   | 2.0007     | 0.9999     | 1.6400      |
| 2   | 2.0032     | 1.0029     | 1.5780      |
| 3   | 2.0025     | 1.0041     | 1.5780      |
| 4   | 1.9846     | 0.9964     | 1.5780      |
| 5   | 1.9975     | 0.9976     | 1.7180      |
| 6   | 2.0027     | 1.0116     | 1.7030      |
| 7   | 1.9970     | 1.0068     | 1.6570      |
| 8   | 1.6450     | 1.0006     | 1.5940      |
| 9   | 2.0003     | 1.0100     | 1.5930      |

Cuadro 4.6: Orden de Convergencia  $factor = 100$  (FEMGM)

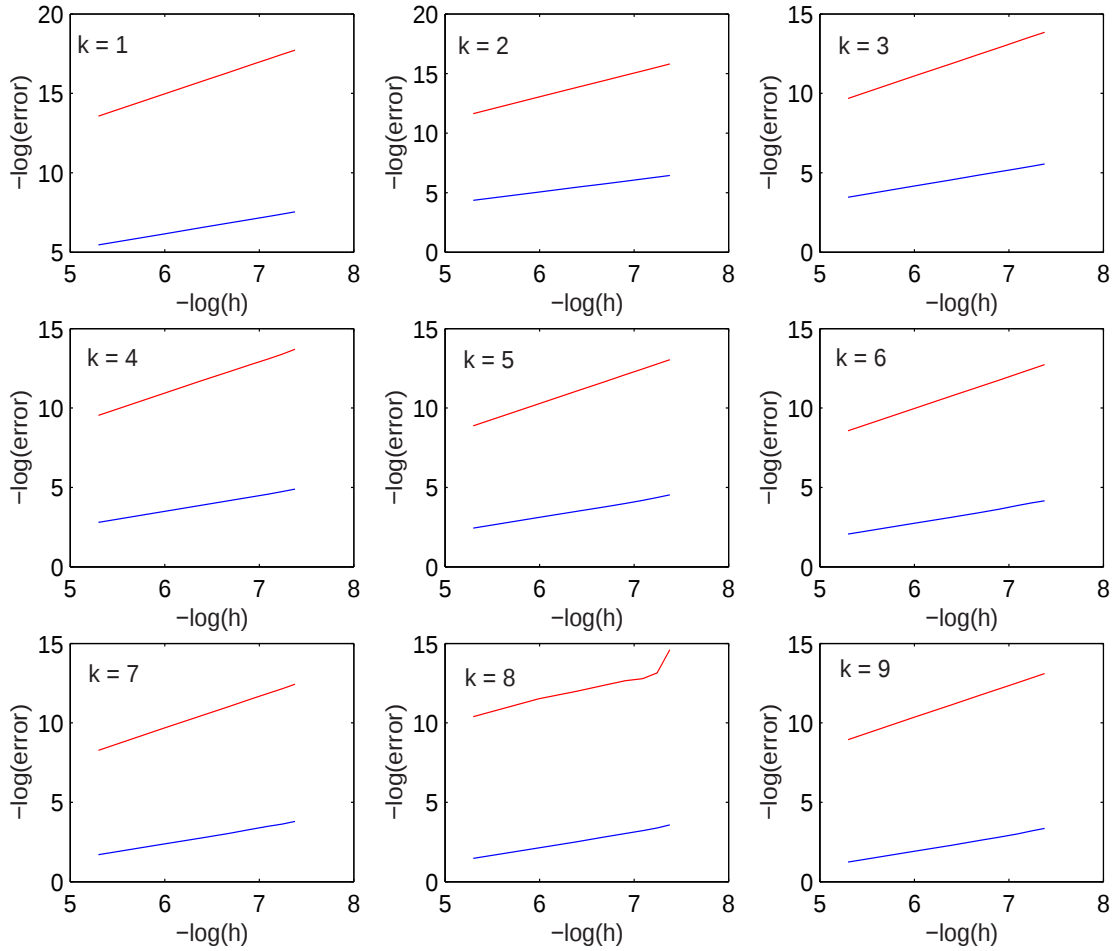


Figura 4.7: Error  $factor = 200$  (FEMGM)

| $k$ | $Orden(p)$ | $Orden(z)$ | $tiempo(s)$ |
|-----|------------|------------|-------------|
| 1   | 2.0004     | 1.0000     | 4.3280      |
| 2   | 2.0015     | 1.0026     | 4.1880      |
| 3   | 2.0012     | 1.0040     | 4.2500      |
| 4   | 1.9838     | 0.9964     | 4.2500      |
| 5   | 1.9986     | 0.9980     | 4.2500      |
| 6   | 2.0008     | 1.0114     | 4.2030      |
| 7   | 1.9954     | 1.0064     | 4.2190      |
| 8   | 1.6436     | 1.0007     | 4.2190      |
| 9   | 1.9995     | 1.0100     | 4.2340      |

Cuadro 4.7: Orden de Convergencia  $factor = 200$  (FEMGM)

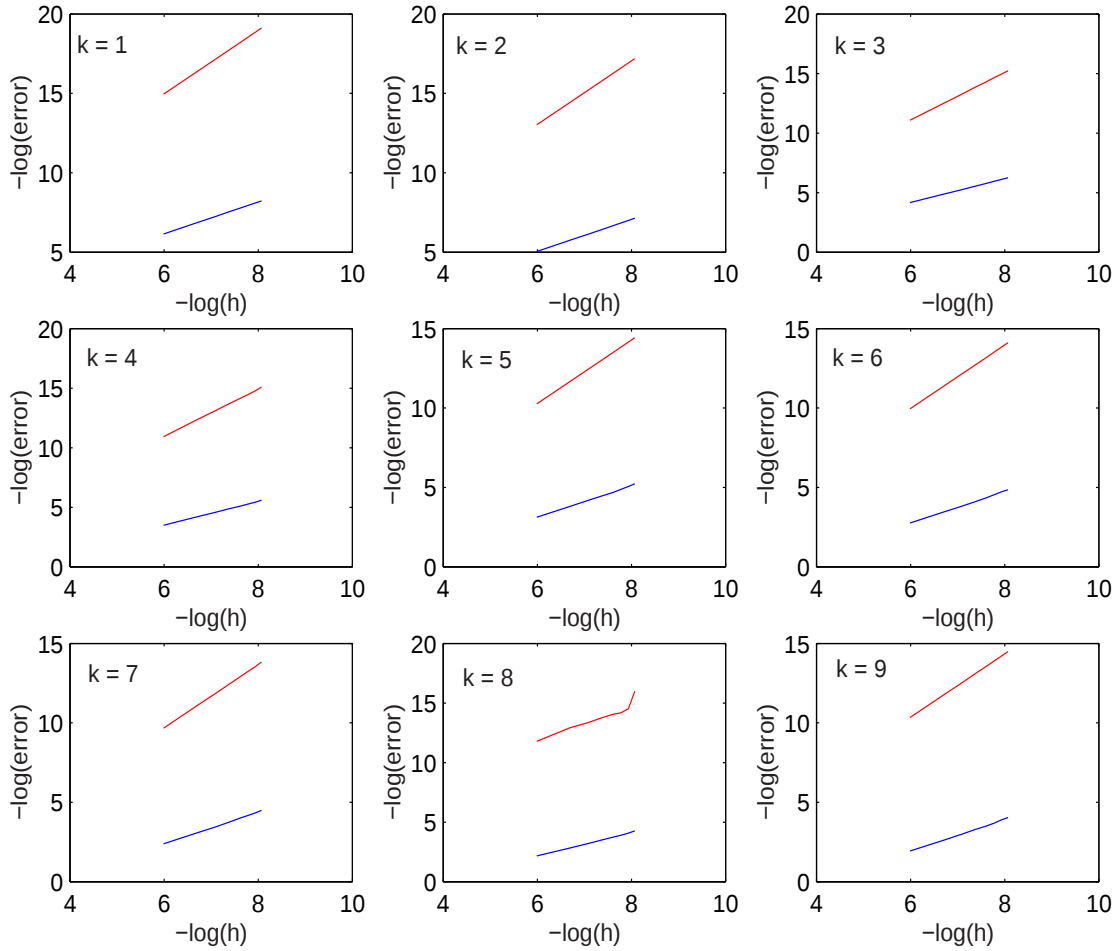


Figura 4.8: Error  $factor = 400$  (FEMGM)

| $k$ | $Orden(p)$ | $Orden(z)$ | $tiempo(s)$ |
|-----|------------|------------|-------------|
| 1   | 2.0002     | 1.0000     | 13.8910     |
| 2   | 2.0006     | 1.0025     | 13.0630     |
| 3   | 2.0005     | 1.0040     | 13.2030     |
| 4   | 1.9835     | 0.9964     | 13.2660     |
| 5   | 1.9989     | 0.9983     | 13.2500     |
| 6   | 1.9998     | 1.0113     | 13.1250     |
| 7   | 1.9947     | 1.0062     | 13.3900     |
| 8   | 1.6429     | 1.0008     | 13.1880     |
| 9   | 1.9993     | 1.0100     | 13.3130     |

Cuadro 4.8: Orden de Convergencia  $factor = 400$  (FEMGM)

### 4.1.3. Método de Elementos Finitos Mínimos Cuadrados para Sistemas de Primer Orden

Como se puede observar en la figura 4.9 para valores de  $k$  menores que seis se tiene un comportamiento lineal para la gráfica de  $-\log(error)$  en función de  $-\log(h)$ , siendo las mismas casi paralelas, también se observa que para valores de  $k$  mayores que seis se tiene un comportamiento de una curva creciente cóncava y a medida que  $-\log(h)$  aumenta tiende a comportarse lineal, esto se evidencia cuando se incrementa el valor de la variable  $factor$  para los mismos valores de  $k$ , obsérvese las figuras 4.10, 4.11 y 4.12. En todos los casos se puede observar en los cuadros 4.9, 4.10, 4.11 y 4.12 que la pendiente de dichas rectas tiende a dos, tanto para la variable primal  $p$  como para la secundaria  $z$ . En el anexo C se presenta el código en *MATLAB* generado para realizar dichos cálculos.

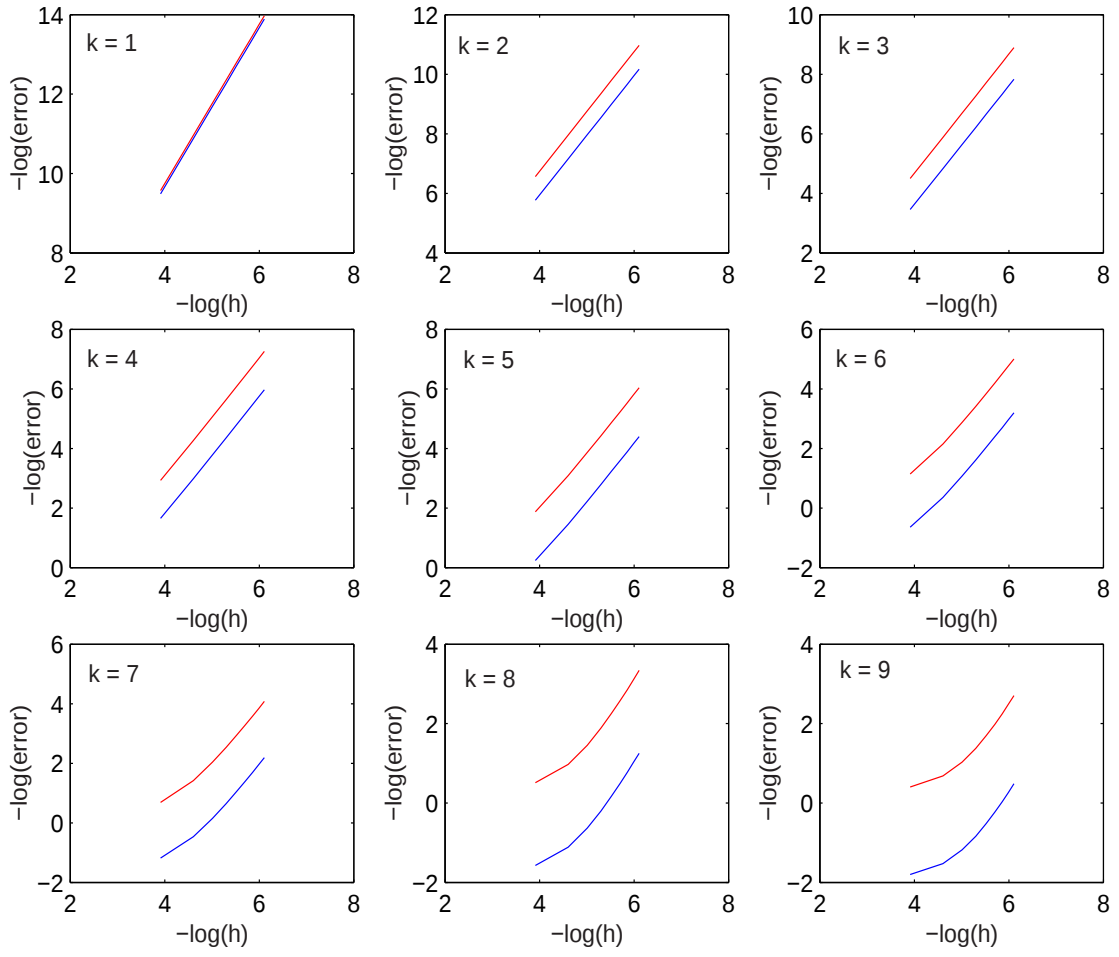


Figura 4.9: Error  $factor = 50$  (LSFEM)

| $k$ | $Orden(p)$ | $Orden(z)$ | $tiempo(s)$ |
|-----|------------|------------|-------------|
| 1   | 2.0038     | 2.0036     | 0.7190      |
| 2   | 2.0058     | 2.0003     | 0.7350      |
| 3   | 2.0000     | 1.9931     | 0.7190      |
| 4   | 1.9736     | 1.9717     | 0.7350      |
| 5   | 1.9090     | 1.9072     | 0.7350      |
| 6   | 1.7873     | 1.7780     | 0.7350      |
| 7   | 1.5850     | 1.5789     | 0.7350      |
| 8   | 1.3340     | 1.3361     | 0.7030      |
| 9   | 1.0883     | 1.0804     | 0.7030      |

Cuadro 4.9: Orden de Convergencia  $factor = 50$  (LSFEM)

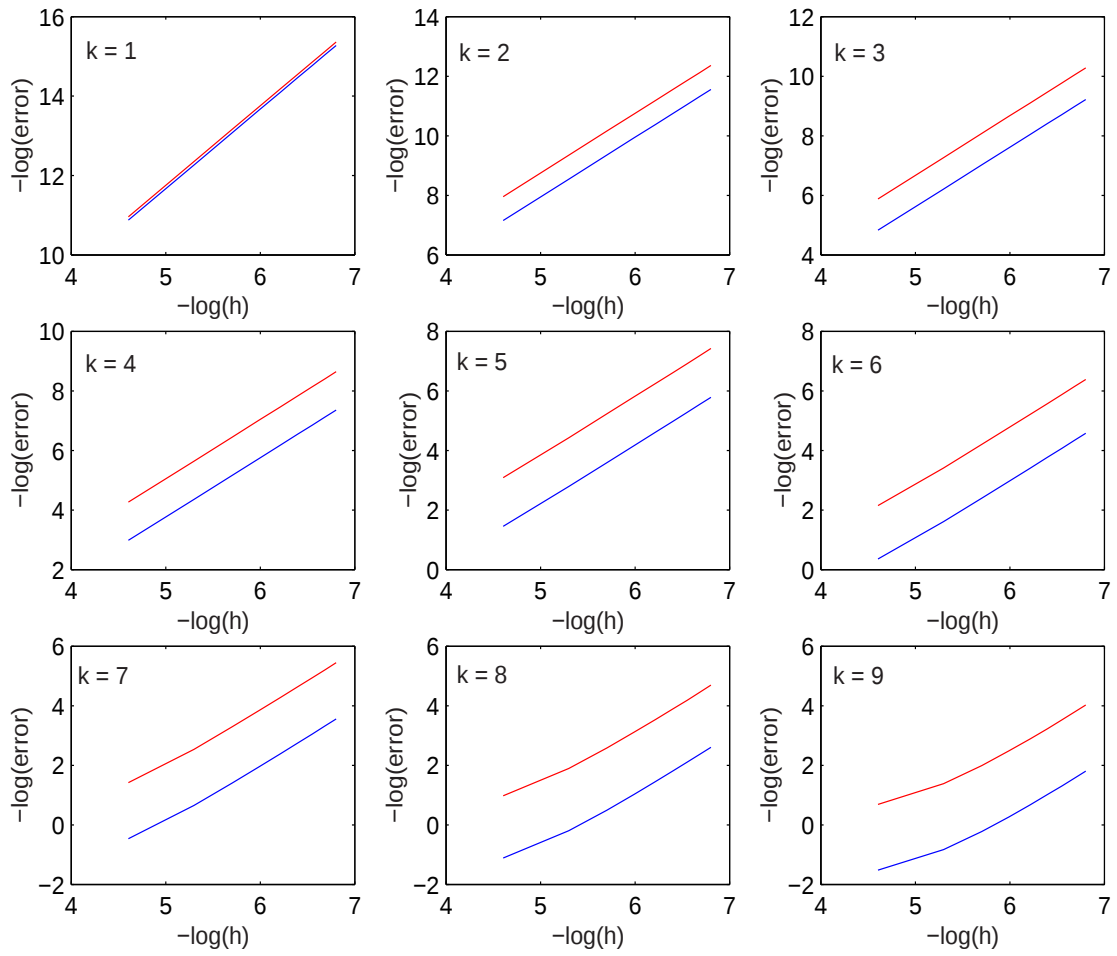


Figura 4.10: Error  $factor = 100$  (LSFEM)

| $k$ | $Orden(p)$ | $Orden(z)$ | $tiempo(s)$ |
|-----|------------|------------|-------------|
| 1   | 2.0016     | 2.0017     | 1.6410      |
| 2   | 2.0032     | 2.0004     | 1.6100      |
| 3   | 2.0016     | 1.9981     | 1.6090      |
| 4   | 1.9936     | 1.9932     | 1.6100      |
| 5   | 1.9755     | 1.9751     | 1.6090      |
| 6   | 1.9371     | 1.9320     | 1.6100      |
| 7   | 1.8541     | 1.8517     | 1.6250      |
| 8   | 1.7231     | 1.7253     | 1.6090      |
| 9   | 1.5619     | 1.5576     | 1.6100      |

Cuadro 4.10: Orden de Convergencia  $factor = 100$  (LSFEM)

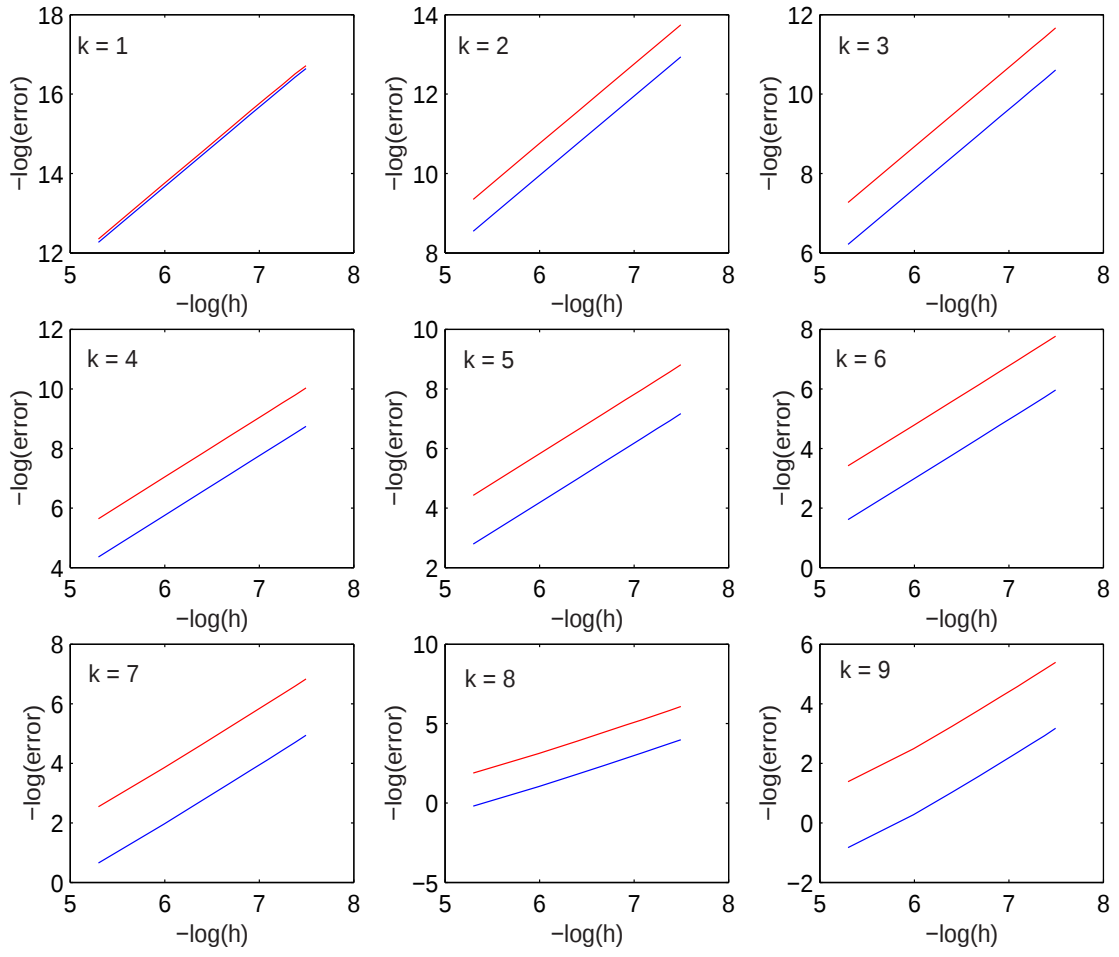


Figura 4.11: Error  $factor = 200$  (LSFEM)

| $k$ | $Orden(p)$ | $Orden(z)$ | $tiempo(s)$ |
|-----|------------|------------|-------------|
| 1   | 1.9936     | 1.9968     | 4.0470      |
| 2   | 2.0013     | 1.9999     | 3.9530      |
| 3   | 2.0011     | 1.9993     | 3.9530      |
| 4   | 1.9982     | 1.9984     | 4.0150      |
| 5   | 1.9933     | 1.9936     | 3.9680      |
| 6   | 1.9838     | 1.9806     | 4.0620      |
| 7   | 1.9577     | 1.9569     | 3.9840      |
| 8   | 1.9108     | 1.9132     | 3.9690      |
| 9   | 1.8436     | 1.8411     | 3.9690      |

Cuadro 4.11: Orden de Convergencia  $factor = 200$  (LSFEM)

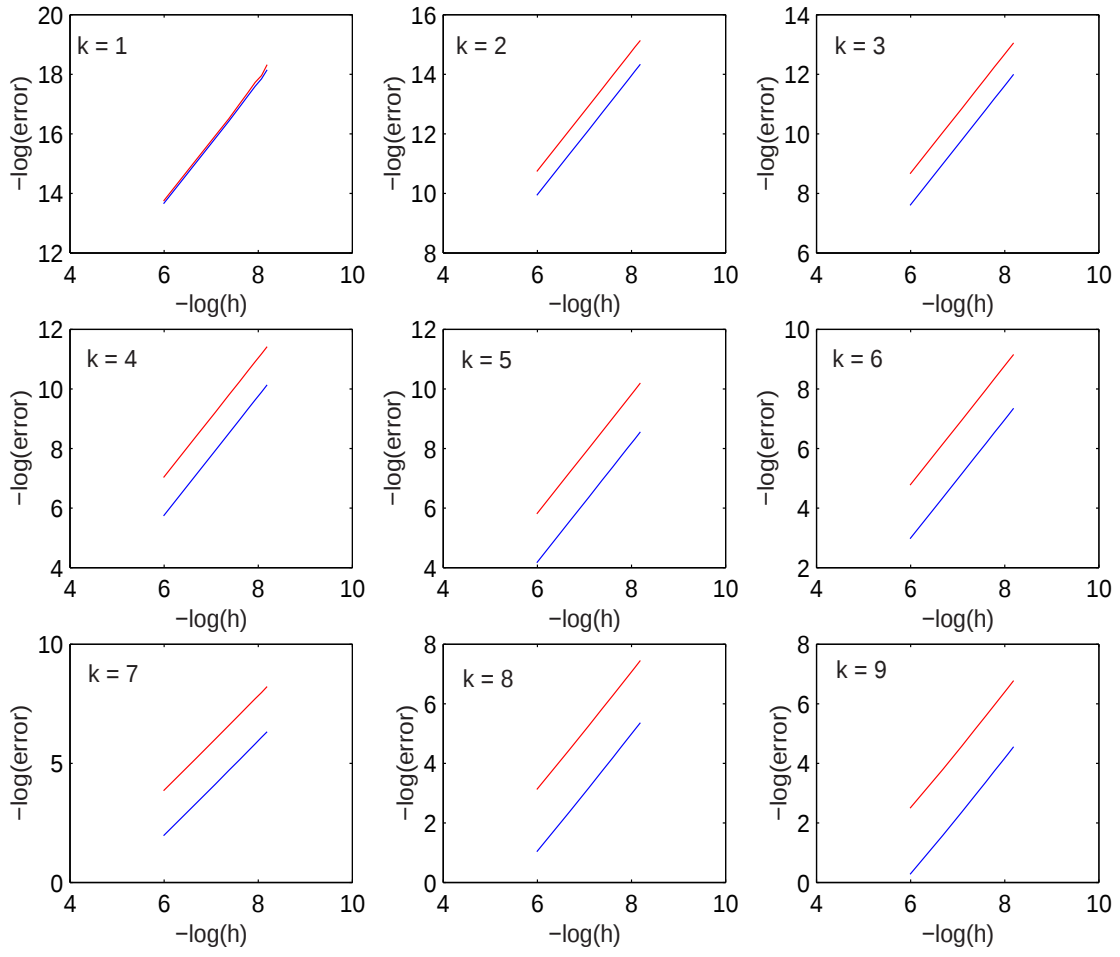


Figura 4.12: Error  $factor = 400$  (LSFEM)

| $k$ | $Orden(p)$ | $Orden(z)$ | $tiempo(s)$ |
|-----|------------|------------|-------------|
| 1   | 2.0638     | 2.0347     | 17.8900     |
| 2   | 2.0031     | 2.0029     | 16.4840     |
| 3   | 2.0013     | 2.0005     | 16.4530     |
| 4   | 1.9994     | 1.9999     | 16.7350     |
| 5   | 1.9977     | 1.9984     | 16.3910     |
| 6   | 1.9960     | 1.9938     | 16.1880     |
| 7   | 1.9879     | 1.9877     | 16.9690     |
| 8   | 1.9734     | 1.9761     | 17.6090     |
| 9   | 1.9541     | 1.9525     | 16.2030     |

Cuadro 4.12: Orden de Convergencia  $factor = 400$  (LSFEM)

## 4.2. Análisis de la Dispersión

### Introducción

Ihlenburg e Ihlenburg y Babuška en [Ih1] y [Ih2] respectivamente, encontraron que el error en la solución de la ecuación de Helmholtz por medio del método de los elementos finitos Galerkin contiene un término de contaminación, el cual depende del número de onda  $k$ ; este término tiene mayor influencia a medida que  $k$  aumenta. En esta sección se hará un análisis de la diferencia entre el número de onda de la solución analítica y el de la solución numérica que no es más que el comúnmente denominado error de contaminación por efecto de la dispersión. El error total es la suma del error de interpolación más el error debido al efecto de la contaminación, esto se observa en la figura 4.13.

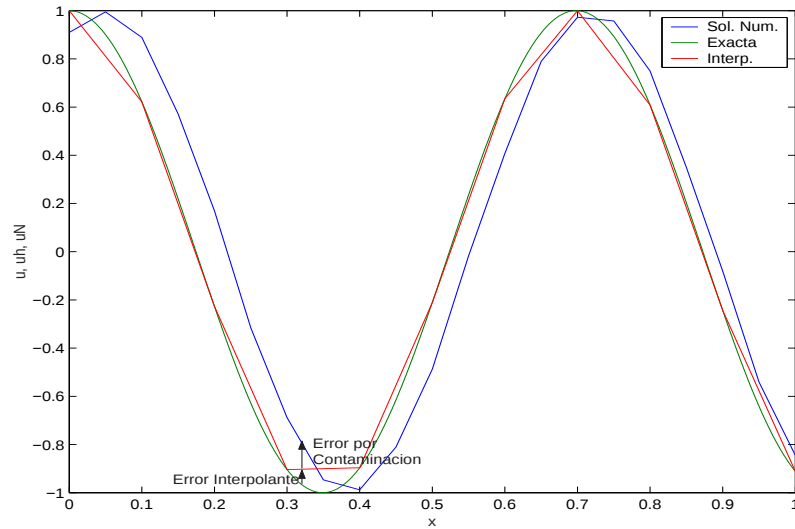


Figura 4.13: Solución exacta, interpolante y solución numérica

Para calcular el número de onda de la solución numérica para los métodos numéricos usados en esta investigación se utilizarán las ecuaciones en diferencias obtenidas

nidas previamente. Así se tiene a lo sumo dos ecuaciones en diferencia con seis incógnitas ( tres en  $z$  y tres en  $p$ ) que para resolverlas se utilizará el procedimiento dado en [Lev] el cual se describirá a continuación.

Sean las dos ecuaciones en diferencias

$$f(n, x_n, x_{n+1}, \dots, x_{n+r}, y_n, y_{n+1}, y_{n+2}) = 0 \quad (4.1)$$

$$g(n, x_n, x_{n+1}, \dots, x_{n+r}, y_n, y_{n+1}, y_{n+2}) = 0 \quad (4.2)$$

Para obtener una sola ecuación en  $x$  y  $n$  se tienen que eliminar  $y_n$ ,  $y_{n+1}$  y  $y_{n+2}$ . Para ello en las ecuaciones (4.1) y (4.2) se sustituye  $n$  por  $n + 1$ , obteniéndose dos nuevas ecuaciones y una nueva incógnita

$$f(n + 1, x_{n+1}, x_{n+2}, \dots, x_{n+r+1}, y_{n+1}, y_{n+2}, y_{n+3}) = 0 \quad (4.3)$$

$$g(n + 1, x_{n+1}, x_{n+2}, \dots, x_{n+r+1}, y_{n+1}, y_{n+2}, y_{n+3}) = 0 \quad (4.4)$$

En este momento se tienen cuatro ecuaciones y se desean eliminar las cuatro variables  $y_n$ ,  $y_{n+1}$ ,  $y_{n+2}$  y  $y_{n+3}$ . Luego se sustituye  $n$  por  $n + 2$  en las ecuaciones (4.1) y (4.2) obteniéndose

$$f(n + 2, x_{n+2}, x_{n+3}, \dots, x_{n+r+2}, y_{n+2}, y_{n+3}, y_{n+4}) = 0 \quad (4.5)$$

$$g(n + 2, x_{n+2}, x_{n+3}, \dots, x_{n+r+2}, y_{n+2}, y_{n+3}, y_{n+4}) = 0 \quad (4.6)$$

De esta forma se tienen seis ecuaciones y cinco variables a eliminar. Ahora si se puede proceder a eliminar estas cinco variables de estas seis ecuaciones realizando simples manipulaciones algebraicas , generándose así una sola ecuación en  $x$  y  $n$  como la siguiente

$$f_1(n, x_n, x_{n+1}, x_{n+2}, \dots, x_{n+r}, x_{n+r+1}, x_{n+r+2}) = 0 \quad (4.7)$$

la cual se resuelve utilizando los procedimientos básicos para obtener la solución de una ecuación en diferencia ordinaria lineal, para ello consúltese el capítulo cuatro *Linear Difference and Functional Equations with Constant Coefficients* del texto [Lev].

Para resolver la ecuación en diferencias ordinaria lineal con coeficientes constantes (4.7) de orden  $r + 3$  se sustituye

$$x_i = X_i \lambda^i \quad (4.8)$$

con lo que se obtiene una ecuación polinómica en  $\lambda$  de grado  $r + 2$  la cual se resuelve obteniendo sus  $r + 2$  raíces. En el caso de esta investigación se pueden obtener expresiones algebraicamente exactas ya que  $r = 2$ , o sea se obtiene una ecuación polinómica de cuarto grado, la cual se resuelve utilizando el procedimiento presentado en [Wei] el cual es debido a Ferrari. El resultado es también verificado utilizando el programa *MAPLE V*. Véase los apéndices (*E*, *G* y *H*).

Las ecuaciones en diferencia obtenidas en los métodos de diferencias finitas implícito (3.3) y (3.4), elemento finito Galerkin Mixto (3.21) y (3.29) así como Mínimos Cuadrados (3.46) y (3.54) se pueden escribir como un sistema de ecuaciones en diferencia como el siguiente

$$Bz_{n-1} + 2Az_n + Bz_{n+1} + Dp_{n-1} - Dp_{n+1} = 0 \quad (4.9)$$

$$-Dz_{n-1} + Dz_{n+1} + Fp_{n-1} + 2Ep_n + Fp_{n+1} = 0 \quad (4.10)$$

donde  $A, B, D, E$  y  $F$  dependen de  $h$  y  $k$  y son diferentes para cada método. A este sistema de ecuaciones en diferencias se le aplica el procedimiento descrito anteriormente para obtener una ecuación polinómica de cuarto grado, la cual se

puede escribir como

$$U\lambda^4 + V\lambda^3 + W\lambda^2 + V\lambda + U = 0 \quad (4.11)$$

donde  $U, V$  y  $W$  dependen de  $A, B, D, E$  y  $F$ . Después se procede a resolver dicha ecuación de cuarto grado, obteniéndose los valores  $\lambda_1 = \bar{\lambda}_2$  y  $\lambda_3 = \bar{\lambda}_4$ , donde  $\bar{z}$  denota el conjugado del número complejo  $z$ . Como es bien sabido  $\lambda = |\lambda|e^{i\tilde{k}h}$  [Ihl], en donde considerando la parte real y despejando  $\tilde{k}$  obtenemos las ecuaciones para calcular  $\tilde{k}_1$  y  $\tilde{k}_2$

$$\tilde{k}_1 = \frac{1}{h} \arccos\left[\frac{Re(\lambda_1)}{|\lambda_1|}\right] \quad (4.12)$$

y

$$\tilde{k}_2 = \frac{1}{h} \arccos\left[\frac{Re(\lambda_3)}{|\lambda_3|}\right] \quad (4.13)$$

A continuación se presentan los resultados de los dos métodos de elementos finitos tratados en esta investigación. La secuencia a seguir en las próximas secciones será primero la identificación de los coeficientes de las ecuaciones en diferencias (4.9) y (4.10), así como los de la ecuación polinómica (4.11). Luego se presentan tabulados los valores de  $\tilde{k}$  para diversos valores de  $h$  y  $k$  después se analizarán todos los elementos que conforman la ecuaciones (4.12) y (4.13), así como también el error porcentual en la aproximación numérica del número de onda.

Otro hecho ocurrido cuando se caracteriza el número de onda es la frecuencia de corte (*Cutoff*) que no es más que la frecuencia normalizada ( $kh$ ) para la cual el valor absoluto de  $\lambda$  cambia bruscamente, bien sea expandiéndose o contrayéndose. En [Ih2] se hace referencia a dicha frecuencia para el método de elementos finitos Galerkin, pero sólo en casos expansivos.

### 4.2.1. Método de Diferencias Finitas Implícito

Los coeficientes de las ecuaciones en diferencias (4.9) y (4.10) para el método de diferencias finitas implícito son  $A = h$ ,  $B = 0$ ,  $D = 1$ ,  $E = hk^2$  y  $F = 0$ , los cuales se obtienen de las ecuaciones (3.3) y (3.4).

Para este método, en el caso de la ecuación polinómica (4.11) se tienen los coeficientes  $U = 1$ ,  $V = 0$  y  $W = 4h^2k^2 - 2$ . Al resolver dicha ecuación polinómica y calculando  $\tilde{k}$  con las ecuaciones (4.12) y (4.13) para diferentes valores de  $h$  y  $k$  se obtienen los valores que se presentan tabulados en el Cuadro 4.13. En el anexo D se presenta el código en *MAPLE* generado para realizar dichos cálculos.

| $h$    | $k = 2$    | $k = 6$    | $k = 9$    |
|--------|------------|------------|------------|
| 0.1    | 2.01357921 | 6.43501109 | 11.1976951 |
| 0.05   | 2.00334842 | 6.09385308 | 9.33530678 |
| 0.01   | 2.00013336 | 6.00360584 | 9.01219450 |
| 0.005  | 2.00003333 | 6.00090036 | 9.00304027 |
| 0.001  | 2.00000133 | 6.00003600 | 9.00012150 |
| 0.0005 | 2.00000033 | 6.00000900 | 9.00003037 |
| 0.0001 | 2.00000001 | 6.00000036 | 9.00000122 |

Cuadro 4.13:  $\tilde{k}$  para diversos valores de  $h$  y  $k$  para el Método de diferencias finitas implícito

A continuación se presenta un estimado para el error relativo en el cálculo del número de onda numérico debido al efecto de dispersión de la solución numérica. Para ello se considerará la solución analítica de la ecuación (4.11) para el caso que

estamos tratando, obteniéndose

$$\lambda_1 = \sqrt{-2(kh)^2 + 1 + 2\sqrt{(kh)^4 - (kh)^2}} \quad (4.14)$$

$$\lambda_2 = -\sqrt{-2(kh)^2 + 1 + 2\sqrt{(kh)^4 - (kh)^2}} \quad (4.15)$$

$$\lambda_3 = \sqrt{-2(kh)^2 + 1 - 2\sqrt{(kh)^4 - (kh)^2}} \quad (4.16)$$

$$\lambda_4 = -\sqrt{-2(kh)^2 + 1 - 2\sqrt{(kh)^4 - (kh)^2}} \quad (4.17)$$

Utilizando la ecuación (4.12) y desarrollando  $\arccos\left[\frac{\operatorname{Re}(\lambda_1)}{|\lambda_1|}\right]$  en serie de Taylor con  $\lambda_1$  dado por (4.14) se tiene que

$$\tilde{k}h = kh + \frac{1}{6}(kh)^3 + \frac{3}{40}(kh)^5 + \frac{89285625}{199999800}(kh)^7 + O((kh)^9) \quad (4.18)$$

Utilizando la definición de error relativo y (4.18) se obtiene

$$E_{r,k} = \frac{\tilde{k} - k}{k} = \frac{\tilde{k}h - kh}{kh} = \frac{1}{6}(kh)^2 + \frac{3}{40}(kh)^4 + O((kh)^6) \quad (4.19)$$

Observe que el error relativo en el cálculo de  $k$  ( $E_{r,k}$ ) es de orden  $O((kh)^2)$ . Un procedimiento similar se utiliza para la ecuaciones (4.15), (4.16) y (4.17) obteniéndose también el resultado dado por (4.19) en los tres casos. Para los detalles de los cálculos analice el anexo D.

En las gráficas de la figura 4.14 se presentan el número de onda numérico en función del número de onda para diferentes niveles de discretización, en ellas se puede observar que para los valores "pequeños" de  $k$  dicha curva se comporta como una recta ( $\tilde{k} = k$ ) que es lo deseado. Luego se observa a partir de un valor de  $k$ , dependiente de  $h$ , que no se comporta de esta manera, incrementándose sustancialmente el error. Véase las gráficas del error porcentual para la aproximación de  $k$  en la figura 4.15

En la figura 4.16 se presentan gráficas de la parte real, parte imaginaria y valor absoluto de los diferentes valores de  $\lambda$  al cual en las gráficas se le denomina  $a$ ,

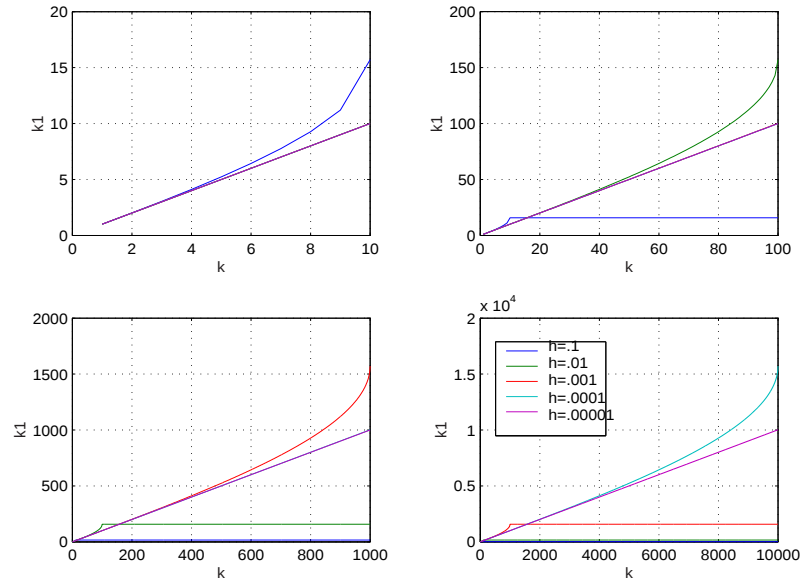


Figura 4.14:  $\tilde{k}$  versus  $k$  para diferentes valores de  $h$  para el método de diferencias finitas implícito

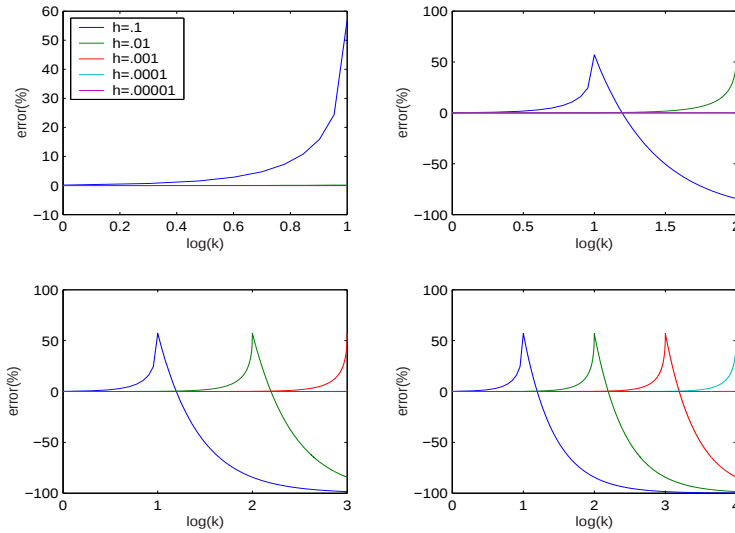


Figura 4.15: Error porcentual en la aproximación de  $\tilde{k}$  para el método de diferencias finitas implícito

observándose un valor de frecuencia de corte igual a la unidad. En estas gráficas se presentan también a las funciones  $\cos(kh)$  y  $\sin(kh)$ , o la negativas de ellas según el caso. En las mismas se observa que para  $\lambda_1$  y  $\lambda_2$  ocurre la contracción

brusca en el valor absoluto de  $\lambda_1$  y  $\lambda_2$  para la frecuencia de corte y para  $\lambda_3$  y  $\lambda_4$  ocurre la expansión.

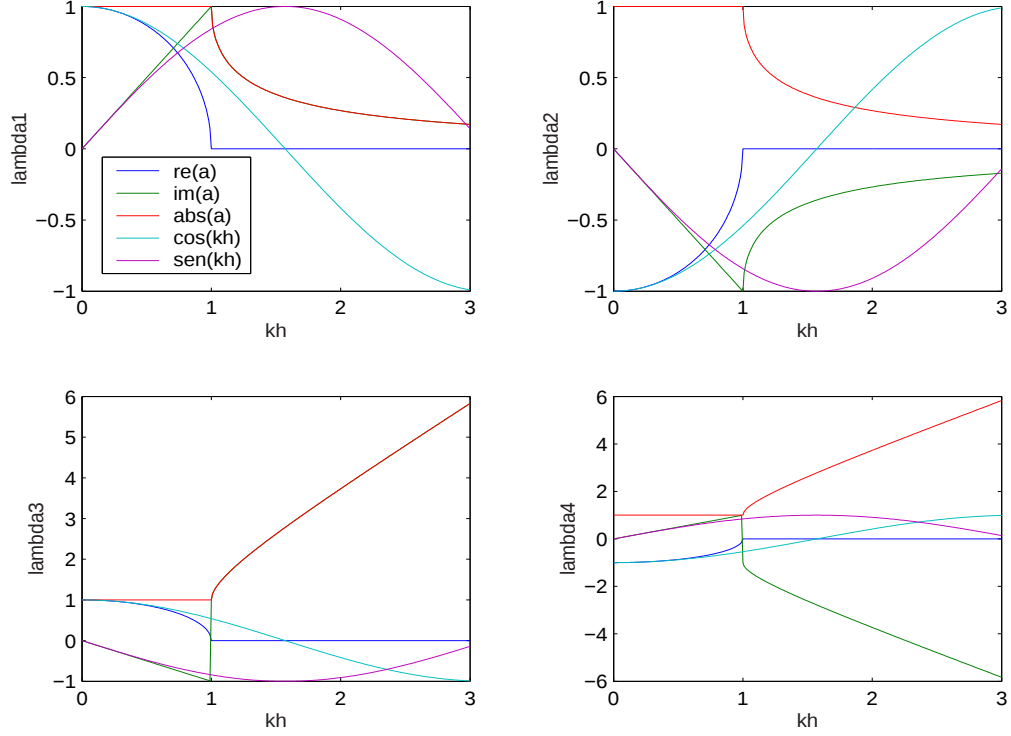


Figura 4.16: Gráficas de la parte real, imaginaria y valor absoluto de los diferentes valores de  $\lambda$  para el método de diferencias finitas implícito

#### 4.2.2. Método de Elementos Finitos Galerkin Mixto

Los coeficientes de las ecuaciones en diferencias (4.9) y (4.10) para el método de elementos finitos Galerkin Mixto son  $A = \frac{h}{3}$ ,  $B = \frac{h}{6}$ ,  $D = \frac{1}{2}$ ,  $E = \frac{h}{3}k^2$  y  $F = \frac{h}{6}k^2$ , los cuales se obtienen de las ecuaciones (3.21) y (3.29).

Para este método en el caso de la ecuación polinómica (4.11) se tienen los coeficientes  $U = BF + D^2$ ,  $V = 2AF + 2BE$  y  $W = 4AE + 2DF - 2D^2$ . Al resolver dicha ecuación polinómica y calculando  $\tilde{k}$  con las ecuaciones (4.12) y (4.13) para

diferentes valores de  $h$  y  $k$  se obtienen los valores que se presentan tabulados en el Cuadro 4.14. En el anexo E se presenta el código en *MAPLE* generado para realizar dichos cálculos.

| $h$    | $k = 2$    | $k = 6$    | $k = 9$    |
|--------|------------|------------|------------|
| 0.1    | 2.00001786 | 6.00452634 | 9.03689143 |
| 0.05   | 2.00000111 | 6.00027297 | 9.00210280 |
| 0.01   | 2.00000000 | 6.00000043 | 9.00000328 |
| 0.005  | 2.00000000 | 6.00000003 | 9.00000021 |
| 0.001  | 2.00000000 | 6.00000000 | 9.00000000 |
| 0.0005 | 2.00000000 | 6.00000000 | 9.00000000 |
| 0.0001 | 2.00000000 | 6.00000000 | 9.00000000 |

Cuadro 4.14:  $\tilde{k}$  para diversos valores de  $h$  y  $k$  para el Método de elementos finitos Galerkin Mixto

Para obtener un estimado para el error relativo en el cálculo del número de onda numérico debido al efecto de dispersión de la solución numérica por medio del método de elementos finitos Galerkin Mixto se consideraran las solución analítica de la ecuación (4.11) para el caso que estamos tratando

$$\lambda_1 = \frac{-2(hk)^2 + 3\sqrt{9 - 3(hk)^2}}{9 + (hk)^2} + \frac{hk\sqrt{3(hk)^2 - 12\sqrt{9 - 3(hk)^2} - 45}}{9 + (hk)^2} \quad (4.20)$$

$$\lambda_2 = \frac{-2(hk)^2 + 3\sqrt{9 - 3(hk)^2}}{9 + (hk)^2} - \frac{hk\sqrt{3(hk)^2 - 12\sqrt{9 - 3(hk)^2} - 45}}{9 + (hk)^2} \quad (4.21)$$

$$\lambda_3 = -\frac{2(hk)^2 + 3\sqrt{9 - 3(hk)^2}}{9 + (hk)^2} + \frac{hk\sqrt{3(hk)^2 + 12\sqrt{9 - 3(hk)^2} - 45}}{9 + (hk)^2} \quad (4.22)$$

$$\lambda_4 = \frac{-2(hk)^2 + 3\sqrt{9 - 3(hk)^2}}{9 + (hk)^2} - \frac{hk\sqrt{3(hk)^2 + 12\sqrt{9 - 3(hk)^2} - 45}}{9 + (hk)^2} \quad (4.23)$$

Utilizando la ecuación (4.12) y desarrollando  $\arccos[\frac{Re(\lambda_1)}{|\lambda_1|}]$  en serie de Taylor con  $\lambda_1$  dado por (4.20) se tiene que

$$\tilde{k}h = kh - \frac{5}{900}(kh)^5 - \frac{5291}{7999992}(kh)^7 + O((kh)^9) \quad (4.24)$$

Utilizando la definición de error relativo y (4.24) se obtiene

$$E_{r,k} = \frac{\tilde{k} - k}{k} = -\frac{5}{900}(kh)^4 - \frac{5291}{7999992}(kh)^6 + O((kh)^8) \quad (4.25)$$

Observe que el error relativo en el cálculo de  $k$  ( $E_{r,k}$ ) es de orden  $O((kh)^4)$ . Para los detalles de los cálculos observe el anexo E.

En las gráficas de las figuras 4.17 y 4.18 se presentan el número de onda numérico en función del número de onda para diferentes niveles de discretización y error porcentual cometido en dicha aproximación respectivamente. Los resultados obtenidos son similares a los obtenidos para el método de diferencias finitas implícito.

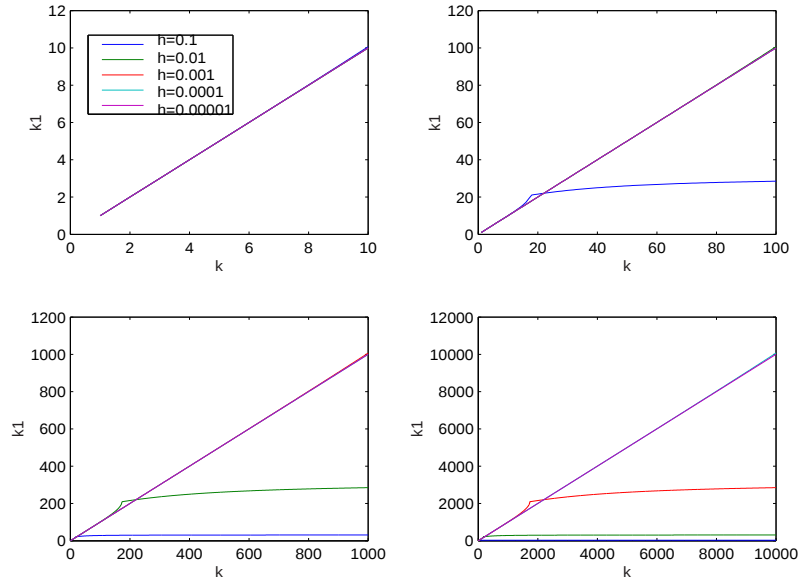


Figura 4.17:  $\tilde{k}$  versus  $k$  para diferentes valores de  $h$  para el método de elementos finitos Galerkin Mixto

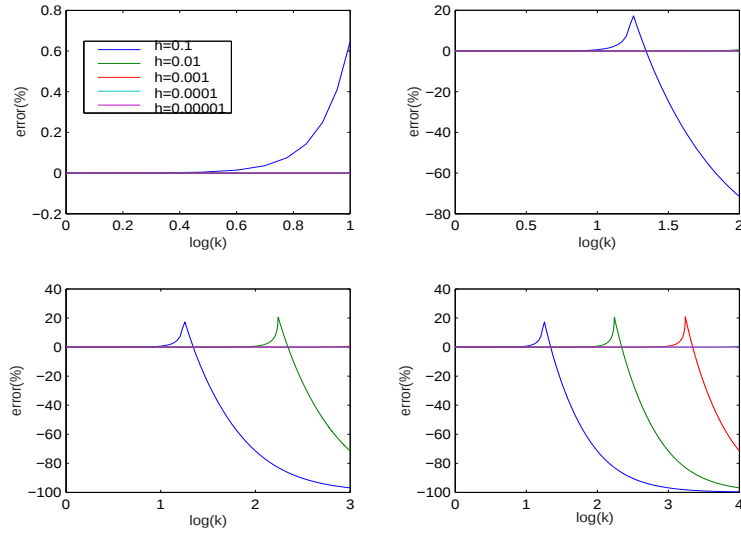


Figura 4.18: Error porcentual en la aproximación de  $k$  para el método de elementos finitos Galerkin Mixto

En la figura 4.19 se presentan las curvas para entender el comportamiento de los diferentes valores de  $\lambda$ , observándose una frecuencia de corte igual a  $\sqrt{3}$ .

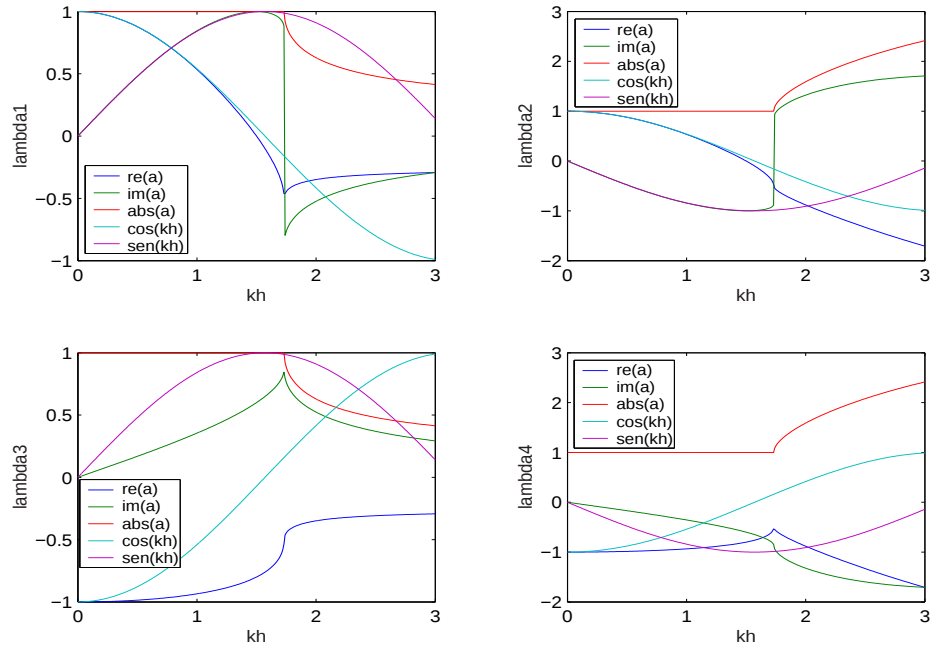


Figura 4.19: Gráficas de la parte real, imaginaria y valor absoluto de los diferentes valores de  $\lambda$  para el método de elementos finitos Galerkin Mixto

### 4.2.3. Método de Elementos Finitos Mínimos Cuadrados para Sistemas de Primer Orden

Los coeficientes de las ecuaciones en diferencias (4.9) y (4.10) para el método de elementos finitos Mínimos Cuadrados para Sistemas de Primer Orden son  $A = \frac{h}{3} + \frac{1}{h}$ ,  $B = \frac{h}{6} - \frac{1}{h}$ ,  $D = \frac{1+k^2}{2}$ ,  $E = \frac{h}{3}k^4 + \frac{1}{h}$  y  $F = \frac{h}{6}k^4 - \frac{1}{h}$ , los cuales se obtienen de las ecuaciones (3.46) y (3.54).

Para este método en el caso de la ecuación polinómica (4.11) se tienen los mismos coeficientes  $U$ ,  $V$  y  $W$  que para el método de elementos finitos Galerkin Mixto. Al resolver dicha ecuación polinómica, se obtiene

$$\lambda_1 = \frac{a + ib + hk\sqrt{c + id}}{e} \quad (4.26)$$

$$\lambda_2 = \frac{a + ib - hk\sqrt{c + id}}{e} \quad (4.27)$$

$$\lambda_3 = \frac{-a + ib + hk\sqrt{c + id}}{e} \quad (4.28)$$

$$\lambda_4 = \frac{-a + ib - hk\sqrt{c + id}}{e} \quad (4.29)$$

donde

$$a = 36 - 2h^4k^4 + 3h^2(1 + k^4) \quad (4.30)$$

$$b = 3\sqrt{3}h^3k^2(1 + k^2) \quad (4.31)$$

$$c = -36e + h^2k^2(104 - 45h^2 - 45h^2k^4 + 3h^4k^4 - 54h^2k^2) \quad (4.32)$$

$$d = 6\sqrt{3}h(1 + k^2)a \quad (4.33)$$

$$e = h^4k^4 + 3h^2 + 3h^2k^4 + 36 + 18h^2k^2 \quad (4.34)$$

y calculando  $\tilde{k}$  con las ecuaciones (4.12) y (4.13) para diferentes valores de  $h$  y  $k$ , se obtienen los valores que se presentan tabulados en el Cuadro 4.15. En el anexo F se presenta el código en *MAPLE* generado para realizar dichos cálculos.

| $h$    | $k = 2$    | $k = 6$    | $k = 9$    |
|--------|------------|------------|------------|
| 0.1    | 1.98141538 | 4.53310957 | 4.69021555 |
| 0.05   | 1.99528419 | 5.45278584 | 6.54728747 |
| 0.01   | 1.99981046 | 5.97366210 | 8.81553538 |
| 0.005  | 1.99995607 | 5.99337235 | 8.95245965 |
| 0.001  | 1.99999810 | 5.99973434 | 8.99807930 |
| 0.0005 | 1.99999953 | 5.99993358 | 8.99951968 |
| 0.0001 | 1.99999998 | 5.99999734 | 8.99998079 |

Cuadro 4.15:  $\tilde{k}$  para diversos valores de  $h$  y  $k$  para el Método de elementos finitos Mínimos Cuadrados

En las gráficas de las figuras 4.20 y 4.21 se presentan el número de onda numérico en función del número de onda para diferentes niveles de discretización y el error porcentual cometido en dicha aproximación respectivamente. Los resultados obtenidos son similares a los obtenidos para el método de diferencias finitas implícito y elementos finitos Galerkin Mixto.

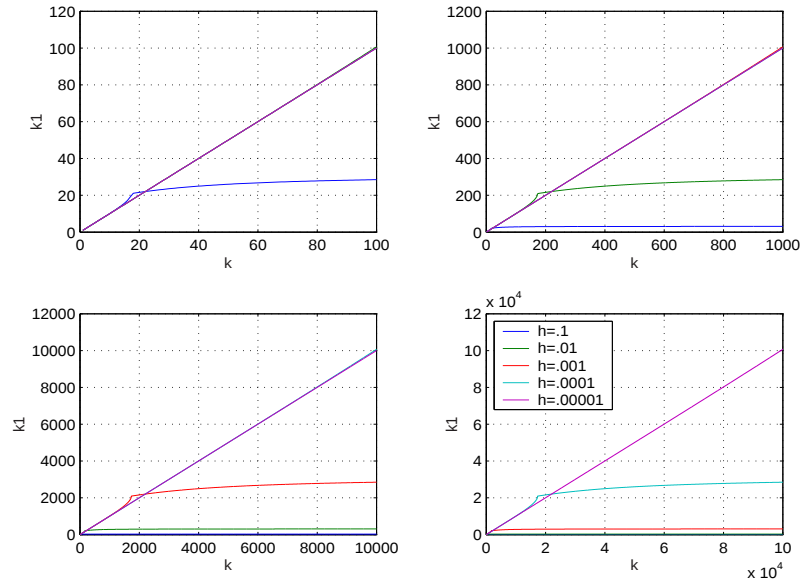


Figura 4.20:  $\tilde{k}$  versus  $k$  para diferentes valores de  $h$  para el Método de elementos finitos Mínimos Cuadrados

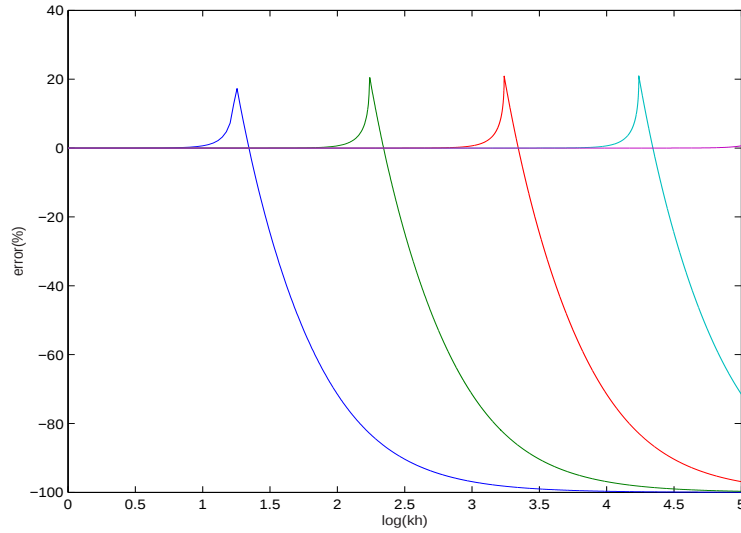


Figura 4.21: Error porcentual en la aproximación de  $k$  para el Método de elementos finitos Mínimos Cuadrados

### 4.3. Comentarios Finales y Conclusiones

Se hace una descripción introductoria a los métodos de diferencia finita y elementos finitos que pudiera utilizarse como referencia en trabajos posteriores y como soporte a las personas que se inician en esta área, ya que se hace una presentación sencilla y debido a la poca cantidad de documentos de este tipo, escritos en este idioma. Sin embargo, en el caso de desear una visión más general y/o profunda de estos temas se recomienda la consulta de material altamente especializados, como por ejemplo los citados en este documento.

Se presentan formulaciones novedosas para los métodos de elementos finitos Galerkin Mixto y mínimos cuadrados que pueden ser utilizados para problemas similares y aún otros diferentes, sin necesidad de modificar sustancialmente dichas formulaciones.

En cuanto a la aplicación al problema de dispersión de onda unidimensional tra-

tado como un sistema de ecuaciones diferenciales ordinarias de primer orden se obtuvo un orden de convergencia igual a dos, para la variable primal  $p$ , en los tres métodos; igual ocurrió con los métodos de diferencias finitas implícito y elementos finitos mínimos cuadrados para la variable secundaria  $z$ , no así para el método de elementos finitos Galerkin Mixto, para el cual se obtuvo un valor de uno.

En el estudio de la contaminación del error se obtuvieron frecuencias de corte igual a uno y  $\sqrt{3}$  para los métodos de diferencias finitas y elementos finitos Galerkin Mixto respectivamente. También se presentan estimados de la contaminación del error para estos métodos.

## 4.4. Trabajos futuros

Formular e implementar los métodos de elementos finitos Galerkin Mixto y Mínimos Cuadrados para sistemas de primer orden con funciones de forma de grado superior, así como métodos de diferencias finitas miméticos de orden superior para el problema de dispersión de ondas unidimensionales.

Hacer comparaciones de estos métodos en cuanto a la contaminación del error debido a la dispersión de la onda numérica.

Hacer un estudio de estabilidad del método de diferencias finitas dominio tiempo propuesto en el anexo H de este trabajo para el problema de dispersión de ondas unidimensional.

Implementar la formulación que se presenta en el anexo J para el problema de dispersión de ondas exterior.

Implementar condiciones DtN, tomando en cuenta las condiciones de irradiación en el infinito de diverso orden para problemas exteriores.

# Apéndice A

## Programación del Método de Diferencias Finitas Implícito

```
factor = 250;          % factor es el valor donde comienza n, el cual se va
% incrementando de factor en factor hasta llegar a
% orden_matriz_graficas al cuadrado por factor.
maximo = 3;          % maximo2 indica la cantidad de recuadros a graficar
% para la convergencia
orden_matriz_graficas = 3;    % para las gráficas de la solución aproximada.
im = sqrt(-1);        % definimos el número complejo i
for k = 1 : 1 : maximo2      % k es la frecuencia
.  t = clock;
.  k;          % esta línea se agregó en caso de querer ver el valor de k (elimine el ;)
.  c = im * k;
.  for n = factor : factor : orden_matriz_graficas2 * factor
.    .  h = 1/n;          % espaciado = discretización uniforme
.    .  a = 2 * k2 * h;
.    .  b = 2 * h;
```

```

. . x = 0 : h : 1;          % nodos
. . kk = sparse(0);         % genera una matriz esparcida con un solo elemento
. . kk(1, 1) = 3;
. . kk(1, 3) = -4;
. . kk(1, 5) = 1;
. . kk(2, 3) = a;
. . kk(2, 4) = 1;
. . for i = 1 : 1 : n - 2    % llenado de kk
. . . j = 2 * i + 1;
. . . kk(j, j - 2) = 1;
. . . kk(j, j - 1) = b;
. . . kk(j, j + 2) = -1;
. . . kk(j + 1, j - 1) = -1;
. . . kk(j + 1, j + 2) = a;
. . . kk(j + 1, j + 3) = 1;
. . end
. . j = 2 * n - 1;
. . kk(j, j - 2) = 1        % elementos que faltan de kk
. . kk(j, j - 1) = b;
. . kk(j, j + 1) = -1/c;
. . kk(j + 1, j - 3) = 1;
. . kk(j + 1, j - 1) = -4;
. . kk(j + 1, j + 1) = 3 + a/c;
. . b1 = zeros(j + 1, 1)    % conformación del término independiente
. . b1(1, 1) = -b * c;
. . b1(2, 1) = c;

```

---

```

%
. . . s = kk1      % Resolución del sistema C*s=b1 (s = solución)
%resto=norm(b1-kk*s,2)
. . . for i = 1 : n
. . . . def(i+1) = real(s(2*i-1));
. . . end
. . . def(1) = 1;
%plot(x,def,'b-');
. . . clear kk;      % libera memoria
. . . for i = 1 : 2 : 2*n-3
. . . . er((i+1)/2,1) = real(s(i)) - real(exp(im*k*x((i+1)/2)));
. . . . erder((i+3)/2,1) = real(s(i+1)) - real(im*k*exp(im*k*x((i+3)/2)));
. . . end
. . . er(2*n,1) = 0;
. . . erder(1,1) = 0;
. . . erro(n/factor,1) = norm(er,2)*sqrt(h);
. . . erroder(n/factor,1) = norm(erder,2)*sqrt(h);
. . . end
. . . q = k;
. . . Elapsed_time(q,1) = etime(clock,t);
. . . figure(1); subplot(maximo,maximo,q);      % Gráfica de la convergencia
. . . z = factor : factor : orden_matriz_graficas^2 * factor;
. . . lerrorU = -log(erro);
. . . lerrorUder = -log(erroder);
. . . lhvector = log(z);
. . . plot(lhvector,lerrorU','r-',lhvector,lerrorUder','b-');

```

```

. xlabel('−log(h)');
. ylabel('−log(error)');
. m(q,1) = k;
. a = reglineal(lhvector', lerrorU);
. m(q,2) = a(1);
. a = reglineal(lhvector', lerrorUder);
. 4m(q,3) = a(1);
end
m
Elapsed_time
clear all

```

## Apéndice B

# Programación del Método de Elementos Finitos Galerkin Mixto

```
% Programa GM para la ecuacion de Helmholtz en 1-D (gm1.m)
% ( $u'' + k^2 * u = 0$ ) con condiciones  $p'(0) = i * k$  y  $p'(1) - i * k * p(1) = 0$ 
clear all;

factor = 100;          % factor es el valor donde comienza  $n$ , el cual se va
.                       % incrementando de factor en factor hasta llegar a
.                       % orden_matriz_graficas al cuadrado por factor.
maximo = 2;            %  $maximo^2$  indica la cantidad de recuadros a
.                       % graficar para la convergencia
orden_matriz_graficas = 2;      % para las gráficas de la solución aproximada.
im = sqrt(-1);          % definimos el número complejo  $i$ 
for k = 1 : 1 : maximo^2      %  $k$  es la frecuencia
.   k          % esta línea se agregó en caso de querer ver el valor de  $K$ (elimine el ;)
.   b = k^2; f = im * k; t = clock;
.   for n = factor : factor : orden_matriz_graficas^2 * factor
```

```

. . h = 1/n;          % espaciamento: discretización uniforme
. . x = 0 : h : 1;    % nodos
. . kk = sparse(0);   % genera una matriz esparcida con un solo elemento
. . a = h/6;
. . c = a * b;        % a, b, c, e, f: variables auxiliares para conformar
. .                  % aux1 y aux2
. . e = 2 * c + f;
. . aux1 = [a, -1/2, 4 * a, 0, a, 1/2];    % con aux1 y aux 2 se forman las
. . aux2 = [-1/2, c, 0, 2 * c, 1/2, c];    % filas de kk (elementos no cero)
. . kk(1,1) = 2 * c; kk(1,2) = 1/2; kk(1,3) = c;    %llenado de kk
. . for i = 1 : n - 1
. . . for j = 1 : 6
. . . . if i == 1
. . . . . if j < 4
. . . . . . kk(1,j) = aux2(j + 3);
. . . . . . if j == 1
. . . . . . . kk(1,1) + = 1/2;
. . . . . . . end
. . . . . . end
. . . . . end
. . . . . end
. . . . . if i == 2
. . . . . . if j < 6
. . . . . . . kk(2,j) = aux1(j + 1);
. . . . . . . kk(3,j) = aux2(j + 1);
. . . . . . . end
. . . . . end
. . . . . end

```

```

. . . . . if  $i \geq 2$ 
. . . . .  $kk(2 * i, j + 2 * i - 3) = aux1(j);$ 
. . . . .  $kk(2 * i + 1, j + 2 * i - 3) = aux2(j);$ 
. . . . . end
. . . end
. . end
. . clear  $aux1, aux2;$ 
. .  $p = 2 * n + 1;$ 
. .  $o = p + 1;$ 
. .  $kk(p, p - 2) = a;$           % elementos que faltan de kk
. .  $kk(p, p - 1) = 1/2;$ 
. .  $kk(p, p) = 2 * a;$ 
. .  $kk(p, p + 1) = -1/2;$ 
. .  $kk(o, o - 3) = -1/2;$ 
. .  $kk(o, o - 2) = c;$ 
. .  $kk(o, o - 1) = -1/2;$ 
. .  $kk(o, o) = e;$ 
. . clear  $a, b;$  clear  $c, e;$           % libera memoria
. .  $b1 = zeros(o, 1);$           % conformación del termino independiente
. .  $b1(1, 1) = 1/2 * f;$ 
. .  $b1(1, 1) = -a * f;$ 
. .  $b1(1, 1) = 1/2 * f;$ 
%-----
. .  $s = kk1;$           % Resolución del sistema  $C*s=b1$  ( $s$  = solución)
. . clear  $b1, kk;$  clear  $p, o;$           % libera memoria
. . if  $factor == n/2$ 

```

```

. . . figure(2);
. . . for j = 1 : n + 1
. . . . y = real(s(2 * j));          % Gráfica de la solución aproximada
. . . end
. . . plot(x, y);
. . end

%-----
. . for i = 1 : 2 : 2 * n + 1
. . . er((i + 1)/2, 1) = real(s(i + 1)) - real(exp(f * x((i + 1)/2)));
. . . erder((i + 1)/2, 1) = real(s(i)) - real(f * exp(f * x((i + 1)/2)));
. . end
. . erro(n/factor, 1) = norm(er, 2) * sqrt(h);
. . erroder(n/factor, 1) = norm(erder, 2) * sqrt(h);
. end
. q = k;
. Elapsed_time(q, 1) = etime(clock, t);
. figure(1); subplot(maximo, maximo, q);          % Gráfica de la convergencia
. z = factor : factor : orden_matriz_graficas*orden_matriz_graficas*factor;
. lerrorU = -log(erro);
. lerrorUder = -log(erroder);
. lhvector = log(z);
. plot(lhvector, lerrorU', 'r-', lhvector, lerrorUder', 'b-');
. xlabel('-log(h)'); ylabel('-log(error)');
. m(q, 1) = k; . a = reglineal(lhvector', lerrorU);
. m(q, 2) = a(1);
. a = reglineal(lhvector', lerrorUder);

```

```
.  m(q,3) = a(1);  
end m Elapsed_time clear all
```

## Apéndice C

# Programación del Método de Elementos Finitos Mínimos Cuadrados para Sistemas de Primer Orden

```
% Programa LSFEM para la ecuación de Helmholtz en 1-D (ls2.m)
% ( $u'' + k^2 * u = 0$ ) con condiciones  $p'(0) = i * k$  y  $p'(1) - i * k * p(1) = 0$ 
clear all;

factor = 1000      ; % factor es el valor donde comienza  $n$ ,
.                  % el cual se va incrementando de factor en factor
.                  % hasta llegar a orden_matriz_graficas al cuadrado por factor.
maximo = 3;        %  $maximo^2$  indica la cantidad de recuadros a graficar
.                  % para la convergencia
orden_matriz_graficas = 3;      % para las gráficas de la solución aproximada.
im = sqrt(-1);      % definimos el número complejo  $i$ 
for k = 1 : 1 : maximo^2      %  $k$  es la frecuencia
.    k                % esta línea se agregó en caso de querer ver el valor de  $k$  (elimine el ;)
```

```

.    $k2 = k^2;$ 
.    $k4 = k^2^2;$ 
.    $t = clock;$ 
.    $for\ n = factor : factor : orden\_matriz\_graficas^2 * factor$ 
.   .    $h = 1/n;$            % espaciamento: discretización uniforme
.   .    $x = 0 : h : 1;$        % nodos
.   .    $kk = sparse(0);$        % genera una matriz esparcida con un solo elemento
.   .    $a = h/3 + 1/h;$ 
.   .    $b = h/6 - 1/h;$ 
.   .    $c = (1 - k2)/2;$ 
.   .    $d = 1 - c;$            %  $a, b, c, d, e, f$ : variables auxiliares
.   .    $e = k4 * h/3 + 1/h;$    % para conformar  $aux1$  y  $aux2$ 
.   .    $f = k4 * h/6 - 1/h;$ 
.   .    $aux1 = [b, d, 2 * a, 0, b, -d];$ 
.   .    $aux2 = [-d, f, 0, 2 * e, d, f];$            % con  $aux1$  y  $aux2$  se forman las
.                                           % filas de  $kk$  (elementos no cero)
.   .    $for\ i = 1 : n - 1$        %llenado de  $kk$ 
.   .   .    $for\ j = 1 : 6$ 
.   .   .   .    $if\ i == 1$ 
.   .   .   .   .    $if\ j < 4$ 
.   .   .   .   .   .    $kk(1, j) = aux2(j + 3);$ 
.   .   .   .   .   .    $if\ j == 1$ 
.   .   .   .   .   .   .    $kk(1, 1) = kk(1, 1)/2;$ 
.   .   .   .   .   .   .    $end$ 
.   .   .   .   .    $end$ 
.   .   .   .    $end$ 

```

```

. . . . . if i == 2
. . . . . if j < 6
. . . . . . kk(2, j) = aux1(j + 1);
. . . . . . kk(3, j) = aux2(j + 1);
. . . . . end
. . . . . end
. . . . . if i >= 2
. . . . . . kk(2 * i, 2 * i + j - 3) = aux1(j);
. . . . . . kk(2 * i + 1, 2 * i + j - 3) = aux2(j);
. . . . . end
. . . . . end
. . . . . end
. . . . . clear aux1, aux2;
. . . . . kk(2 * n, 2 * n - 2) = b ; % elementos que faltan de kk
. . . . . kk(2 * n, 2 * n - 1) = d;
. . . . . kk(2 * n, 2 * n) = a + 1;
. . . . . kk(2 * n, 2 * n + 1) = -c;
. . . . . kk(2 * n + 1, 2 * n - 2) = -d;
. . . . . kk(2 * n + 1, 2 * n - 1) = f;
. . . . . kk(2 * n + 1, 2 * n) = -c;
. . . . . kk(2 * n + 1, 2 * n + 1) = e + k2;
. . . . . clear a, f; clear e; % libera memoria
. . . . . B = sparse(2 * n + 1, 2 * n + 1); % matriz esparcida de orden 2n + 1
. . . . . B(2 * n, 2 * n + 1) = k;
. . . . . B(2 * n + 1, 2 * n) = -k;
. . . . . C = [kkB; B'kk]; % conformación de la matriz C

```

```

. . clear kk, B;          % libera memoria
. . b1 = zeros(4 * n + 2, 1);      % conformación del termino independiente
. . b1(2 * n + 2, 1) = -c * k;
. . b1(2 * n + 3, 1) = -b * k;
. . b1(2 * n + 4, 1) = d * k;
. . clear b, c; clear d;
%-----
. . s = C\1;          % Resolución del sistema  $C * s = b1$  ( $s$  = solución)
. . clear b1, C;      % libera memoria
%-----
. . for i = 1 : 2 : 2 * n + 1
. . . er((i + 1)/2, 1) = s(i) - real(exp(im * k * x((i + 1)/2)));
. . . if i == (2 * n + 1)
. . . . erder((i + 3)/2, 1) = s(i + 1) - real(im * k * exp(im * k * x((i + 3)/2)));
. . . end
. . end
. . erder(1, 1) = 0;
. . erro(n/factor, 1) = norm(er, 2) * sqrt(h);
. . erroder(n/factor, 1) = norm(erder, 2) * sqrt(h);
. end
. q = k;
. Elapsed_time(q, 1) = etime(clock, t);
. figure(1); subplot(maximo, maximo, q);          % Gráfica de la convergencia
. z = factor : factor : orden_matriz_graficas*orden_matriz_graficas*factor;
. lerrorU = -log(erro);
. lerrorUder = -log(erroder);

```

```

.   lhvector = log(z);
.   plot(lhvector, lerrorU', 'r-', lhvector, lerrorUder', 'b-');
.   xlabel('-log(h)');
.   ylabel('-log(error)');
.   m(q, 1) = k;
.   a = reglineal(lhvector', lerrorU);
m(q, 2) = a(1);
.   a = reglineal(lhvector', lerrorUder);
.   m(q, 3) = a(1);
end
m Elapsed_time clear all

```

## Apéndice D

# Cálculo de la Dispersión para el Método de Diferencias Finitas Implícito

### D.1. Ecuación en Diferencias

En esta sección se obtiene una ecuación en diferencias en la variable  $p$  a partir de las ecuaciones (3.3) y (3.4). Para ello se decrementa e incrementa el índice en la ecuación (3.3), obteniéndose

$$2hz_{i-1} + p_{i-2} - p_i = 0 \quad (\text{D.1})$$

$$2hz_{i+1} + p_i - p_{i+2} = 0 \quad (\text{D.2})$$

restando las ecuaciones (D.2) y (D.1)

$$2h(z_{i+1} - z_{i-1}) - p_{i-2} + 2p_i - p_{i+2} = 0 \quad (\text{D.3})$$

al despejar  $z_{i+1} - z_{i-1}$  de la ecuación (3.4) y sustituyendo en la ecuación (D.3) se obtiene

$$p_{i-2} + (4h^2k^2 - 2)p_i + p_{i+2} = 0 \quad (\text{D.4})$$

## D.2. Programa en Maple Para Resolver la Ecuación en Diferencias

```

> restart : Digits := 10 :
> _EnvExplicit := true :
> ec := u * x^4 + v * x^3 + w * x^2 + v * x + u = 0;
> u := 1; v := 0; w := 4 * h^2 * k^2 - 2;
> solve(ec, x);
> x1 := sqrt(-2 * k^2 + 1 + 2 * sqrt(k^4 - k^2));
> x2 := -sqrt(-2 * k^2 + 1 + 2 * sqrt(k^4 - k^2));
> x3 := sqrt(-2 * k^2 + 1 - 2 * sqrt(k^4 - k^2));
> x4 := -sqrt(-2 * k^2 + 1 - 2 * sqrt(k^4 - k^2));
> erro1 := evalf(taylor(arccos(x1 - I * k), k = 0, 20));
> erro2 := evalf(taylor(arccos(x2 + I * k), k = 0, 20));
> erro3 := evalf(taylor(arccos(x3 + I * k), k = 0, 20));
> erro4 := simplify(evalf(Pi) - evalf(taylor(arccos(x4 - I * k), k = 0, 20)));
> error1_3 := expand((,1666666667 * k^3 + ,7500000000e - 1 * k^5 + ,4464285714e -
1 * k^7 + ,3038194444e - 1 * k^9 + ,2237215909e - 1 * k^11 + ,1735276442e - 1 * k^13 +
,1396484375e - 1 * k^15 + ,1155180090e - 1 * k^17)/k) + O(k^18);
> error2_4 := -expand((- ,1666666667 * k^3 - ,7500000000e - 1 * k^5 -
,4464285714e - 1 * k^7 - ,3038194444e - 1 * k^9 - ,2237215909e - 1 * k^11 -
,1735276442e - 1 * k^13 - ,1396484375e - 1 * k^15 - ,1155180090e - 1 * k^17)/k) + O(k^18);

```

## Apéndice E

# Cálculo de la Dispersión para el Método de Elementos Finitos Galerkin Mixto

### E.1. Programa en Maple Para Obtener la Ecuación en Diferencias

```
> restart : with(polytools) :  
> eq1 := b * zj-1 + 2 * a * zj + b * zj+1 + d * pj-1 - d * pj+1 = 0;  
> eq2 := -d * zj-1 + d * zj+1 + f * pj-1 + 2 * e * pj + f * pj+1 = 0 :  
> eq3 := b * zj + 2 * a * zj+1 + b * zj+2 + d * pj - d * pj+2 = 0;  
> eq4 := -d * zj + d * zj+2 + f * pj + 2 * e * pj+1 + f * pj+2 = 0;  
> eq5 := b * zj+1 + 2 * a * zj+2 + b * zj+3 + d * pj+1 - d * pj+3 = 0 :  
> eq6 := -d * zj+1 + d * zj+3 + f * pj+1 + 2 * e * pj+2 + f * pj+3 = 0 :  
> eq7 := d * eq1 + b * eq2 :  
> eq7 := simplify(%);  
> eq8 := d * eq5 - b * eq6 :
```

```

> eq8 := simplify(%);
> zj := (d * zj+2 + f * pj + 2 * e * pj+1 + f * pj+2)/d;
> eq3 := d * eq3 : simplify(%);
> eq7 : simplify(%);
> eq8 : simplify(%);
> zj1 := -(2*b*zj+2*d+b*f*pj+2*b*e*pj+1+b*f*pj+2+d^2*pj-d^2*pj+2)/(2*a*d);
> eq7 := a * eq7 : simplify(%); eq8 := a * eq8 : simplify(%);
> eq9 := (eq7 - eq8)/a : simplify(%);
> def := (b * f + d^2) * pj+1 + (2 * a * f + 2 * b * e) * pj + (4 * a * e + 2 * b * f - 2 *
d^2) * pj+1 + (2 * b * e + 2 * a * f) * pj+2 + (b * f + d^2) * pj+3 = 0;

```

## E.2. Programa en Maple Para Resolver la Ecuación en Diferencias

```

> restart : Digits := 50 :
> _EnvExplicit := true :
> ec := u * x^4 + v * x^3 + w * x^2 + v * x + u = 0;
> u := b * f + d^2; v := 2 * a * f + 2 * b * e; w := 4 * a * e + 2 * b * f - 2 * d^2;
> a := h/3 : b := h/6 : d := 1/2 : e := h/3*k^2 : f := h/6*k^2 : h := 0,0005; k := 2;
> solve(ec, x);
> abs(-,999999994444443672839360425208023731065776039630115-
.
,33333335185185442386879858263992089644741320789962e-3*I);
> abs(,99999950000004166665972222251157392939813609181496-
.
,99999983333334722221990740750385797646604536616054e-3*I);
> k[1] := arccos(-,99999994444443672839360425208023731065776039630115)/h;
> k[2] := arccos(,99999950000004166665972222251157392939813609181496)/h;

```

```
> x1 := (-2 * k^2 + 3 * sqrt(-3 * k^2 + 9))/(k^2 + 9);
> x2 := (-2 * k^2 + 3 * sqrt(-3 * k^2 + 9))/(k^2 + 9) :
> x3 := -(2 * k^2 + 3 * sqrt(-3 * k^2 + 9))/(k^2 + 9);
> x4 := -(2 * k^2 + 3 * sqrt(-3 * k^2 + 9))/(k^2 + 9) :
> erro1 := evalf(taylor(arccos(x1), k = 0, 20));
> error_relativo :=
-expand((,555555555555555555555555555555555555555555555555555555555555556e - 
2 * k^5 + ,66137566137566137566137566137566137566137566137566137566137566137566e - 3 * 
k^7 + ,19290123456790123456790123456790123456790123456790123456790e - 3 * k^9 + 
,43841189674523007856341189674523007856341189674523e - 4 * k^11 + 
,11747190566635011079455523899968344412788857233302e - 4 * k^13 + 
,31257144490169181527206218564243255601280292638317e - 5 * k^15 + 
,86734198208890323390444426513175000224140868052754e - 6 * k^17)/k) +
O(k * 18);
> erro1 := evalf(taylor(arccos(x3), k = 0, 20));
```

## Apéndice F

# Cálculo de la Dispersión para el Método de Elementos Finitos Mínimos Cuadrados para Sistemas de Primer Orden

### F.1. Programa en Maple Para Obtener la Ecuación en Diferencias

```
> restart : with(polytools) :  
> eq1 := b * zj-1 + 2 * a * zj + b * zj+1 + d * pj-1 - d * pj+1 = 0;  
> eq2 := -d * zj-1 + d * zj+1 + f * pj-1 + 2 * e * pj + f * pj+1 = 0 :  
> eq3 := b * zj + 2 * a * zj+1 + b * zj+2 + d * pj - d * pj+2 = 0;  
> eq4 := -d * zj + d * zj+2 + f * pj + 2 * e * pj+1 + f * pj+2 = 0;  
> eq5 := b * zj+1 + 2 * a * zj+2 + b * zj+3 + d * pj+1 - d * pj+3 = 0 :  
> eq6 := -d * zj+1 + d * zj+3 + f * pj+1 + 2 * e * pj+2 + f * pj+3 = 0 :  
> eq7 := d * eq1 + b * eq2 :  
> eq7 := simplify(%);
```

```

> eq8 := d * eq5 - b * eq6 :
> eq8 := simplify(%);
> zj := (d * zj+2 + f * pj + 2 * e * pj+1 + f * pj+2)/d;
> eq3 := d * eq3 : simplify(%);
> eq7 : simplify(%);
> eq8 : simplify(%);
> zj1 := -(2*b*zj+2*d+b*f*pj+2*b*e*pj+1+b*f*pj+2+d^2*pj-d^2*pj+2)/(2*a*d);
> eq7 := a * eq7 : simplify(%); eq8 := a * eq8 : simplify(%);
> eq9 := (eq7 - eq8)/a : simplify(%);
> def := (b * f + d^2) * pj+1 + (2 * a * f + 2 * b * e) * pj + (4 * a * e + 2 * b * f - 2 *
d^2) * pj+1 + (2 * b * e + 2 * a * f) * pj+2 + (b * f + d^2) * pj+3 = 0;

```

## F.2. Programa en Maple Para Resolver la Ecuación en Diferencias. Versión I

```

> restart :
> polinomio := x^4 + a * x^3 + b * x^2 + a * x + 1;

```

Haciendo el cambio de variables

```

> x := z - a/4;

```

Se elimina el término cúbico

```

> simplify(polinomio);
> polinomio := z^4 + (b - 3/8 * a^2) * z^2 + (-1/2 * b * a + a + 1/8 * a^3) * z - 1/4 *
a^2 - 3/256 * a^4 + 1/16 * b * a^2 + 1;
> polinomio := z^4 + p * z^2 + q * z + r;
> p := b - 3/8 * a^2;
> q := -1/2 * b * a + a + 1/8 * a^3;

```

$$> r := -1/4 * a^2 - 3/256 * a^4 + 1/16 * b * a^2 + 1;$$

Hay que resolver la ecuación de tercer grado

$$> ecuacion3 := 4 * (u^2/4 - r) * (u - p) = q;$$

$$> expand(\%, u);$$

$$> solve(u^3 - u^2 * b + 3/8 * u^2 * a^2 + a^2 * u - b * a^2 + 3/8 * a^4 + 3/64 * a^4 * u - 9/64 * a^4 * b + 9/512 * a^6 - 1/4 * b * a^2 * u + 1/4 * b^2 * a^2 - 4 * u + 4 * b - 3/2 * a^2 = -1/2 * b * a + a + 1/8 * a^3, u);$$

$$\begin{aligned} > u_1 := 1/24 * (-1728 * b^2 * a^2 + 864 * a^4 * b - 108 * a^6 - 18432 * b + 6912 * a^2 + 4608 * b * \\ & a^2 - 1728 * a^4 - 3456 * b * a + 6912 * a + 864 * a^3 + 512 * b^3 + 12 * \sqrt{-3145728 + 10368 * \\ & a^8 - 82944 * a^5 - 288768 * a^4 * b^2 + 491520 * b^3 * a^2 + 884736 * b^2 * a - 110592 * b * a^3 + \\ & 49152 * a^2 * b^4 + 26880 * b^4 * a^4 - 21504 * b^3 * a^6 + 7776 * b^2 * a^8 - 122880 * b^3 * a^4 + 96768 * \\ & b^2 * a^6 + 89088 * b^3 * a^3 - 387072 * b^2 * a^3 - 62208 * b^2 * a^5 - 12288 * b^5 * a^2 - 1296 * a^{10} * b - \\ & 27648 * a^8 * b + 221184 * a^5 * b + 15552 * a^7 * b - 24576 * b^4 * a + 49152 * a * b^3 - 196608 * \\ & b^4 + 81 * a^{12} + 2592 * a^{10} - 31104 * a^7 - 1296 * a^9 + 2691072 * a^2 + 663552 * a^3 - 175104 * \\ & a^4 - 2101248 * b * a^2 - 1769472 * b * a + 843264 * a^4 * b - 703488 * b^2 * a^2 - 111552 * a^6 + \\ & 1572864 * b^2))^{(1/3)} - 24 * (-4/3 + 1/3 * a^2 - 1/9 * b^2) / (-1728 * b^2 * a^2 + 864 * a^4 * b - \\ & 108 * a^6 - 18432 * b + 6912 * a^2 + 4608 * b * a^2 - 1728 * a^4 - 3456 * b * a + 6912 * a + 864 * \\ & a^3 + 512 * b^3 + 12 * \sqrt{-3145728 + 10368 * a^8 - 82944 * a^5 - 288768 * a^4 * b^2 + 491520 * \\ & b^3 * a^2 + 884736 * b^2 * a - 110592 * b * a^3 + 49152 * a^2 * b^4 + 26880 * b^4 * a^4 - 21504 * b^3 * \\ & a^6 + 7776 * b^2 * a^8 - 122880 * b^3 * a^4 + 96768 * b^2 * a^6 + 89088 * b^3 * a^3 - 387072 * b^2 * a^3 - \\ & 62208 * b^2 * a^5 - 12288 * b^5 * a^2 - 1296 * a^{10} * b - 27648 * a^8 * b + 221184 * a^5 * b + 15552 * \\ & a^7 * b - 24576 * b^4 * a + 49152 * a * b^3 - 196608 * b^4 + 81 * a^{12} + 2592 * a^{10} - 31104 * a^7 - \\ & 1296 * a^9 + 2691072 * a^2 + 663552 * a^3 - 175104 * a^4 - 2101248 * b * a^2 - 1769472 * b * a + \\ & 843264 * a^4 * b - 703488 * b^2 * a^2 - 111552 * a^6 + 1572864 * b^2))^{(1/3)} + 1/3 * b - 1/8 * a^2; \\ > solve(x^2 + u_1/2 + \sqrt{u_1 - p} * x + q/(2 * \sqrt{u_1 - p}), x); \\ > solve(x^2 + u_1/2 + \sqrt{u_1 - p} * x - q/(2 * \sqrt{u_1 - p}), x); \end{aligned}$$

### F.3. Programa en Maple Para Resolver la Ecuación en Diferencias. Versión II

```
> restart : Digits := 50 : _EnvExplicit := true :  
> ec := u * x^4 + v * x^3 + w * x^2 + v * x + u = 0;  
> u := b * f + d^2; v := 2 * a * f + 2 * b * e; w := 4 * a * e + 2 * b * f - 2 * d^2;  
> a := h/3 + 1/h : b := h/6 - 1/h : c := (1 - k^2)/2 : d := (1 + k^2)/2 :  
> e := h/3 * k^4 + 1/h : f := h/6 * k^4 - 1/h : h := 0,1; k := 9; solve(ec, x);  
> abs(,65877540905580786083954073448283621588186437197659 -  
,33382494972387435069400925814776245900376077269815 * I);  
> abs(1,2078222520924326798219206911502148477987013442922 -  
,61204652911683900391369224228855316698911271887882 * I);  
> k[1] := arccos(,65877540905580786083954073448283621588186437197659/  
,73852835872076991735044376660373578073535924841130)/h;  
> k[2] := arccos(1,2078222520924326798219206911502148477987013442922/  
1,3540441449426992938310411229029442144845099942013)/h;
```

## Apéndice G

### Cálculo de las Integrales Involucradas en los Métodos

A continuación se presentan los cálculos de las integrales involucradas en los métodos de elementos finitos Galerkin Mixto y Mínimos Cuadrados para sistema de primer orden. Para ello se consideran elementos lineales en una discretización uniforme del dominio con funciones de forma definidas por (2.17).

$$\int_0^1 \varphi_0^2 dx = \int_0^h \frac{(x-h)^2}{h^2} dx = \frac{h}{3} \quad (\text{G.1})$$

$$\int_0^1 \varphi_0' \varphi_0 dx = \int_0^h \frac{(x-h)}{h^2} dx = -\frac{1}{2} \quad (\text{G.2})$$

$$\int_0^1 (\varphi_0')^2 dx = \int_0^h \frac{1}{h^2} dx = h \quad (\text{G.3})$$

$$\int_0^1 \varphi_{j-1} \varphi_j dx = \int_{(j-1)h}^{jh} -\frac{(x-(j-1)h)(x-jh)}{h^2} dx = \frac{h}{6} \quad (\text{G.4})$$

$$\int_0^1 \varphi_j^2 dx = \int_{(j-1)h}^{jh} \frac{(x-(j-1)h)^2}{h^2} dx + \int_{jh}^{(j+1)h} \frac{(x-(j+1)h)^2}{h^2} dx = \frac{2h}{3} \quad (\text{G.5})$$

$$\int_0^1 \varphi'_{j-1} \varphi_j dx = \int_{(j-1)h}^{jh} -\frac{(x - (j-1)h)}{h^2} dx = -\frac{1}{2} \quad (\text{G.6})$$

$$\int_0^1 \varphi'_j \varphi_j dx = \int_{(j-1)h}^{jh} \frac{(x - (j-1)h)}{h^2} dx + \int_{jh}^{(j+1)h} -\frac{(x - (j+1)h)}{h^2} dx = 0 \quad (\text{G.7})$$

$$\int_0^1 \varphi'_{j+1} \varphi_j dx = \int_{jh}^{(j+1)h} -\frac{(x - (j+1)h)}{h^2} dx = \frac{1}{2} \quad (\text{G.8})$$

$$\int_0^1 \varphi_n^2 dx = \int_{(n-1)h}^1 \frac{(x - (n-1)h)^2}{h^2} dx = \frac{h}{3} \quad (\text{G.9})$$

$$\int_0^1 \varphi'_n \varphi_n dx = \int_{(n-1)h}^1 \frac{(x - (n-1)h)}{h^2} dx = \frac{1}{2} \quad (\text{G.10})$$

$$\int_0^1 (\varphi'_n)^2 dx = \int_{(n-1)h}^1 \frac{1}{h^2} dx = h \quad (\text{G.11})$$

# Apéndice H

## Método de Diferencias Finitas Dominio Tiempo Explícito

### H.1. Formulación del Método de Diferencias Finitas Dominio Tiempo Explícito

En este caso, a diferencia de la sección anterior, el operador diferencial  $L$  dependerá también del tiempo. Por ello, se utilizan aproximaciones tanto espaciales como temporales para dicho operador. Debido a esto, es necesario establecer una discretización del tiempo, además de la discretización del dominio  $\Omega$  y tener condiciones iniciales. Al discretizar el operador diferencial  $L$ , de la forma ya mencionada, se obtiene la ecuación en diferencias  $L_d^n u_i^n = f_i$ , donde el superíndice  $n$  hace referencia al nivel de tiempo, en el que se ha llevado a cabo la aproximación del operador diferencial  $L$ .

En el caso que en dicha ecuación en diferencias sólo se tenga una incógnita, ya que los demás valores de  $u$  en los nodos y niveles de tiempo involucrados sean conocidos se dirá que se está en presencia de un esquema explícito. Al hacer variar

$i$  en cada nivel de tiempo se obtendrá así todos los valores de  $u$  en los nodos  $x_i$  para el nivel de tiempo  $n$ , está claro que se deben sustituir cuando aparezcan en las ecuaciones en diferencias tanto las condiciones de borde como las iniciales . Luego esto se repite hasta donde se desee incrementando el tiempo en una cantidad  $\Delta t_n$ ;  $n = 1, 2, 3, \dots$ . Usualmente esta variación temporal se establece constante y se le denomina  $\Delta t$ .

## H.2. Método de Diferencias Finitas Dominio Tiempo Explícito

A diferencia de otros métodos presentados en este trabajo, el método de diferencias finitas presentado en esta sección es una formulación que depende del tiempo, por lo que se utilizará la ecuación de onda dispersa, en vez de la ecuación de Helmholtz. El problema de frontera para la onda dispersa se define como

$$(w_{sc})_{tt} = c^2(w_{sc})_{xx}, \quad 0 < x < 1 \quad (\text{H.1})$$

con condición de frontera

$$(w_{sc})_x(0, t) = -(w_{sc})_x(0, t) = ike^{-i\omega x} \big|_{x=0} = ike^{-i\omega t} \quad (\text{H.2})$$

y condición de irradiación en el infinito

$$(w_{sc})_t(1, t) + c(w_{sc})_x(1, t) = 0 \quad (\text{H.3})$$

La expresión para la amplitud de la onda total se descompone como  $w(x, t) = w_{sc}(x, t) + w_{inc}(x, t)$ , donde la onda incidente se expresa como  $w_{inc}(x, t) = e^{-ikx}e^{-i\omega t}$

Es bien conocido que al resolver la ecuación (H.1) sujeta a las condiciones ya establecidas se obtiene por solución  $(w_{sc})(x, t) = e^{ikx}e^{-i\omega t}$

Como la solución del problema de la ecuación de onda es armónica en el tiempo, no importa cuales son las condiciones iniciales, ya que lo que nos interese es la etapa estacionaria. Por ello se pueden utilizar como condiciones iniciales  $(w_{sc})(x, 0) = 0$  y  $(w_{sc})_t(x, 0) = 0$

Para resolver el problema numéricamente y compararlo con el método de elementos finitos mínimos cuadrados se planteará un sistema de ecuaciones diferenciales de primer orden equivalente, a saber:

$$\frac{\partial w}{\partial t} + c \frac{\partial w}{\partial x} = v \quad (\text{H.4})$$

$$\frac{\partial v}{\partial t} - c \frac{\partial v}{\partial x} = 0 \quad (\text{H.5})$$

$$(w_{sc})_x(0, t) = ike^{-ikt} \quad (\text{H.6})$$

$$v(1, t) = 0 \quad (\text{H.7})$$

$$(w_{sc})(x, 0) = 0, (w_{sc})_t(x, 0) = 0 \quad (\text{H.8})$$

En la ecuación (H.6) aparece la frecuencia espacial  $k$  en vez de la frecuencia temporal  $\omega$ , debido a que se ha considerado en la ecuación que las relaciona  $c = \frac{\omega}{k}$  un valor de  $c = 1$ .

Discretizando este sistema de ecuaciones diferenciales utilizando aproximaciones de las derivadas en base a las fórmulas de diferencias centradas de tres puntos dadas por (2.7), tenemos que:

$$\frac{w_i^{n+1} - w_i^{n-1}}{2\Delta t} + \frac{w_{i+1}^n - w_{i-1}^n}{2\Delta x} = v_i^n \quad (\text{H.9})$$

$$\frac{v_i^{n+1} - v_i^{n-1}}{2\Delta t} + \frac{v_{i+1}^n - v_{i-1}^n}{2\Delta x} = 0 \quad (\text{H.10})$$

donde  $i = 1, 2, \dots, N_1 - 1$  para (H.9) e  $i = 2, 3, \dots, N_1 - 1$  para (H.10)

Haciendo lo mismo para las condiciones de frontera, se obtiene

$$\frac{w_2^n - w_0^n}{2\Delta x} = ik e^{-ikt_n} \quad (\text{H.11})$$

$$v_{N_1}^n = 0 \quad (\text{H.12})$$

$$w_i^0 = 0 \quad (\text{H.13})$$

$$\frac{w_i^1 - w_i^{-1}}{2\Delta t} = 0 \quad (\text{H.14})$$

Esta última ecuación implica que  $w_i^{-1} = w_i^1$ , y como ya se ha mencionado que se esta trabajando con un problema cuya solución es armónica en el tiempo se pueden seleccionar  $w_i^{-1} = w_i^1 = 0$ .

Las ecuaciones en diferencias a utilizar para este método se obtienen sustituyendo diferentes valores de  $i$  en las ecuaciones anteriores.

Despejando  $w_0^n$  de (H.11) y sustituyendo  $i = 1$  en (H.9), se obtienen

$$w_0^n = w_2^n - 2ik\Delta x e^{-ikt_n} \quad (\text{H.15})$$

$$w_1^{n+1} = w_1^{n-1} - \frac{\Delta t}{\Delta x}(w_2^n - w_0^n) + 2\Delta t v_1^n \quad (\text{H.16})$$

Sustituyendo (H.15) en (H.16) y luego simplificando, se obtiene

$$w_1^{n+1} = w_1^{n-1} - 2ik\Delta t e^{ikt_n} + 2\Delta t v_1^n \quad (\text{H.17})$$

Luego de las ecuaciones (H.9) y (H.10) para  $i = 2, \dots, N_1 - 1$ , se obtiene

$$w_i^{n+1} = w_i^{n-1} - \frac{\Delta t}{\Delta x}(w_{i+1}^n - w_{i-1}^n) + 2\Delta t v_i^n \quad (\text{H.18})$$

$$v_i^{n+1} = v_i^{n-1} + \frac{\Delta t}{\Delta x}(v_{i+1}^n - v_{i-1}^n) \quad (\text{H.19})$$

La siguiente condición se usará como Criterio de Parada

$$\max_{1 \leq i \leq N_1} \{|w_i^{n+1}| - |w_i^n|, |v_i^{n+1}| - |v_i^n|\} < \varepsilon \quad (\text{H.20})$$

Para hacer la comparación con los otros métodos desarrollados en este trabajo hay que tomar en cuenta que

$$w(x, t) = \widehat{w}(x)e^{-i\omega t} \quad (\text{H.21})$$

$$-i\omega\widehat{w}(x)e^{-i\omega t} + \widehat{w}_x(x)e^{-i\omega t} = \widehat{v}(x)e^{-i\omega t} \quad (\text{H.22})$$

En esta última ecuación despejamos la variable que vamos a comparar con  $z$  en la formulación *LSFEM*

$$\widehat{w}_x(x) = \widehat{v}(x) + i\omega\widehat{w}(x) \quad (\text{H.23})$$

donde

$$\widehat{v}(x) = v(x, t)e^{i\omega t} \quad (\text{H.24})$$

# Apéndice I

## Programación del Método de Diferencias Finitas Dominio Tiempo Explícito

```
% Este programa resuelve la ecuación de Helmholtz basándose en la ecuación de
% onda para ello se resuelve esta última (marchando en el tiempo) hasta que
% se estabilize, esta versión utiliza solo tres filas para p y dos para z
clear all;      % Estos son los parámetros que se le van a pasar a la función
k = 1 * 3,141592654
N1 = 100;
tolerancia = 0,001;
maximo_iteraciones = 300000;
delta_x = 1/(N1 - 1);      % Aquí comienza la función time_domain
delta_t = delta_x/10;
im = sqrt(-1);
P(1, 2 : N1 + 1) = ones(1, N1);
P(2, 2 : N1 + 1) = 1 - im * k * delta_t;
error = 1;
n = 1;
```

```

P(1,1) = P(1,3) - 2 * im * k * delta_x;
P(2,1) = P(2,3) - 2 * im * k * delta_x * exp(-im * k * delta_t);
virtual = P(1, N1) - 2 * delta_x/delta_t * (P(2, N1 + 1) - P(1, N1 + 1));
Z(1, 1 : N1 + 1) = (P(2, 1 : N1 + 1) - P(1, 1 : N1 + 1))/delta_t - (virtual - P(1, 1 :
N1 + 1))/delta_x;
while and(n < maximo_iteraciones, error > tolerancia)
.   Z(2, 2 : N1 + 1) = Z(1, 2 : N1 + 1)
.       - delta_t/delta_x * (Z(1, 2 : N1 + 1) - Z(1, 1 : N1));
.   P(3, 2 : N1) = P(2, 2 : N1)
.       + delta_t/delta_x * (P(2, 3 : N1 + 1) - P(2, 2 : N1))
.       + delta_t * Z(2, 2 : N1);
.   P(3, N1 + 1) = ((3 - delta_t/delta_x) * P(2, N1 + 1)
.       + delta_t/delta_x * P(2, N1)
.       + delta_t * Z(2, N1 + 1))/3;
.   if n == maximo_iteraciones - 1
.       .   P(3,1) = P(3,3) - 2 * im * k * delta_x * exp(-im * k * (n + 1) * delta_t);
.       .   Z(2,1) = (P(3,1) - P(2,1))/delta_t - (P(2,2) - P(2,1))/delta_x;
.       end
.   error1 = max(abs(abs(P(3, 2 : N1 + 1)) - abs(P(2, 2 : N1 + 1))));
.   error2 = max(abs(abs(Z(2, 2 : N1 + 1)) - abs(Z(1, 2 : N1 + 1))));
.   error = max(error1, error2);
.   n = n + 1;
.   for j = 1 : 50
.       .   if n == j * 10000
.           .       n
.           end
.       end

```

```

.   end

.    $P(1, 1 : N1 + 1) = P(2, 1 : N1 + 1);$ 

.    $P(2, 1 : N1 + 1) = P(3, 1 : N1 + 1);$ 

.   if and( $n < maximo\_iteraciones$ ,  $error > tolerancia$ )

.       .    $Z(1, 1 : N1 + 1) = Z(2, 1 : N1 + 1);$ 

.   end

end

N1;n;error1;error2;error

aux1 = real( $P(1, 2 : N1 + 1) * exp(im * k * (n - 1) * delta\_t)$ );
aux2 = real( $P(2, 2 : N1 + 1) * exp(im * k * (n) * delta\_t)$ );
%aux3=real( $-(im*k*P(1,2:N1+1)+Z(1,2:N1+1))*exp(im*k*(n-2)*delta\_t)$ );
aux4 = real( $-(im * k * P(1, 2 : N1 + 1) + Z(2, 2 : N1 + 1)) * exp(im * k * (n -$ 
1) *  $delta\_t)$ );

z = 0 : delta_x : 1;

figure(1);

plot(z, aux1, 'r', z, aux2, 'b');

figure(2);

plot(z, aux4, 'b');

%y=0:delta_x:1+delta_x;

figure(3);

A =  $P(1 : 2, 2 : N1 + 1);$ 

plot(z, A, 'r');

figure(4);

B =  $Z(1 : 2, 2 : N1 + 1);$ 

plot(z, B, 'b');

clearall

```

## **Apéndice J**

# **Formulación y Aplicación del Método de Elemento Finito Mínimos Cuadrados a un Problema de Dispersión de Onda en 2D**

### **Resumen**

En este trabajo se desarrollan dos formulaciones numéricas del método de elemento finito del tipo mínimos cuadrados para resolver la ecuación de Helmholtz en dos dimensiones (2D). Se hacen comparaciones de dichos métodos en cuanto a orden de convergencia, error por contaminación y dispersión. Se presentan varios ejemplos que validan las conclusiones.

## J.1. Ecuación de Helmholtz en 2D

### J.1.1. Introducción

El fenómeno de dispersión de una onda plana que choca con un cilindro infinito con cavidad vacía o hueca (Figura 1) se modela por la ecuación de onda

$$\frac{\partial^2 \bar{p}}{\partial t^2} - c^2 \Delta \bar{p} = 0 \quad (\text{J.1})$$

Como la onda  $\bar{p}$  es armónica en el tiempo ella genera una respuesta armónica en el tiempo (PAL), así

$$\bar{p}(x, t) = p(x)e^{-i\omega t} \quad (\text{J.2})$$

Al sustituir (J.2) en (J.1) se obtiene la ecuación denominada ecuación de onda reducida o ecuación de Helmholtz

$$\Delta p + \frac{\omega^2}{c^2} p = 0 \quad (\text{J.3})$$

donde  $\omega$  es la frecuencia de una onda sinusoidal particular (en tiempo) y  $c$  es la velocidad del sonido. La relación  $k \equiv \frac{\omega}{c}$  es denominada "número de onda".

Consideraremos el caso bidimensional en coordenadas polares

$$\frac{1}{r} \frac{\partial}{\partial r} \left( r \frac{\partial p}{\partial r} \right) + \frac{1}{r^2} \frac{\partial^2 p}{\partial \theta^2} + k^2 p = 0 \quad (\text{J.4})$$

sujeta a la condición de frontera de cavidad hueca  $p(a) = -p_{inc(a)}$ . Para que dicho problema tenga solución única hace falta una condición. Se utilizará una condición de *Non Reflecting Boundary Condition (NRBC)* o condición de frontera absorbente y viene dada por la ecuación

$$\frac{\partial p}{\partial r} - ikp + \frac{p}{2r} \rightarrow_{r \rightarrow \infty} 0 \quad (\text{J.5})$$

## J.2. Método de Elemento Finito del Tipo Mínimos Cuadrados para Sistemas de Primer Orden

Al igual que en una dimensión se desea plantear el sistema de ecuaciones lineales dado por

$$Re(Au, Av) + Re\langle Bu, Bv \rangle_\gamma = 0 \quad (J.6)$$

Antes de proceder a desarrollar dicho sistema vamos a encontrar un sistema equivalente a la ecuación de Helmholtz 2D que sea elíptico. Esto se hará de dos formas diferentes. La primera, se basará en un cambio de variables en función del gradiente y la segunda, en un cambio de variables que involucra a la condición de irradiación.

### J.2.1. Primer Sistema

Para trabajar esta parte reescribiremos la ecuación de Helmholtz de la siguiente forma

$$\frac{\partial p}{\partial r} + r \frac{\partial^2 p}{\partial r^2} + \frac{1}{r} \frac{\partial^2 p}{\partial \theta^2} + k^2 r p = 0 \quad (J.7)$$

Utilizando el cambio de variables  $\nabla p = \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} \frac{\partial p}{\partial r} \\ \frac{1}{r} \frac{\partial p}{\partial \theta} \end{bmatrix}$  se obtiene el sistema de ecuaciones diferenciales de primer orden

$$u + r \frac{\partial u}{\partial r} + \frac{\partial v}{\partial \theta} + k^2 r p = 0 \quad (J.8)$$

$$u - \frac{\partial p}{\partial r} = 0 \quad (J.9)$$

$$rv - \frac{\partial p}{\partial \theta} = 0 \quad (J.10)$$

Es bien sabido que un sistema de ecuaciones diferenciales de tres ecuaciones con tres incógnitas nunca será elíptico. Por ello vamos a agregar una nueva ecuación diferencial de primer orden, combinando las ecuaciones (J.9) y (J.10), a esta nueva ecuación le daremos el nombre de ecuación de compatibilidad. Para obtener la ecuación de compatibilidad derivamos (J.9) respecto a  $\theta$  y (J.10) respecto a  $r$ , la ecuación resultante de esta última se la restamos a la de la primera y se obtiene

$$v + r \frac{\partial v}{\partial r} - \frac{\partial u}{\partial \theta} = 0 \quad (\text{J.11})$$

Ahora plantearemos un sistema de ecuaciones que es equivalente al sistema (J.8)-(J.11). Luego se demostrará que son equivalentes y después se demostrará la elipticidad del último sistema.

Sea el sistema

$$u + r \frac{\partial u}{\partial r} + \frac{\partial v}{\partial \theta} + k^2 r p = 0 \quad (\text{J.12})$$

$$u + \frac{\partial \vartheta}{\partial \theta} - \frac{\partial p}{\partial r} = 0 \quad (\text{J.13})$$

$$r v - r^2 \frac{\partial \vartheta}{\partial r} - \frac{\partial p}{\partial \theta} = 0 \quad (\text{J.14})$$

$$v + r \frac{\partial v}{\partial r} - \frac{\partial u}{\partial \theta} - r \frac{\partial \vartheta}{\partial r} = 0 \quad (\text{J.15})$$

Para demostrar que este sistema es equivalente al anterior basta que  $\vartheta = \text{Constante}$ , procederemos a probar esto

Primero, derivamos las ecuaciones (J.13) respecto a  $\theta$  y (J.14) respecto a  $r$ . Luego a esta última le restamos la anterior, obteniéndose

$$v + r \frac{\partial v}{\partial r} - \frac{\partial u}{\partial \theta} - r \frac{\partial \vartheta}{\partial r} - r^2 \frac{\partial^2 \vartheta}{\partial r^2} - r \frac{\partial \vartheta}{\partial r} - \frac{\partial^2 \vartheta}{\partial \theta^2} = 0 \quad (\text{J.16})$$

Utilizando (J.15) obtenemos

$$-r^2 \frac{\partial^2 \vartheta}{\partial r^2} - r \frac{\partial \vartheta}{\partial r} - \frac{\partial^2 \vartheta}{\partial \theta^2} = 0 \quad (\text{J.17})$$

o lo que es lo mismo  $\Delta\vartheta = 0$ . Si se utiliza la condición de frontera  $\frac{\partial\vartheta}{\partial\eta} = 0$ , entonces se concluye que  $\vartheta = \text{Constante}$ . Por ello los sistemas (J.8-J.11) y (J.12-J.15) son equivalentes.

Ahora comprobaremos que el sistema (J.12-J.15) es elíptico y así (J.8-J.11) también lo es, para ello lo escribimos de la forma  $Au = (A_2\frac{\partial}{\partial\theta} + A_1\frac{\partial}{\partial r} + A_0)u$  donde

$$A_2 = \begin{bmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ -1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \end{bmatrix} \text{ y } A_1 = \begin{bmatrix} 0 & r & 0 & 0 \\ -1 & 0 & 0 & 0 \\ 0 & 0 & 0 & -r^2 \\ 0 & 0 & r & -r \end{bmatrix}$$

Para demostrar la elipticidad de un sistema de ecuaciones diferenciales de primer orden expresado en la forma  $Au = (A_2\frac{\partial}{\partial\theta} + A_1\frac{\partial}{\partial r} + A_0)u$  hay que probar que

$$\det(A_2\eta + A_1\varepsilon) \neq 0 \quad \forall \begin{bmatrix} \varepsilon \\ \eta \end{bmatrix} \neq \begin{bmatrix} 0 \\ 0 \end{bmatrix}$$

$$\text{En nuestro caso tenemos que } \det(A_2\eta + A_1\varepsilon) = (\eta^2 + r^2 \varepsilon^2)^2 \neq 0 \quad \forall \begin{bmatrix} \varepsilon \\ \eta \end{bmatrix} \neq \begin{bmatrix} 0 \\ 0 \end{bmatrix}$$

Así el sistema (J.12)-(J.15) es elíptico

### J.2.2. Segundo Sistema

Para trabajar esta parte reescribiremos la ecuación de Helmholtz de la siguiente forma

$$\frac{\partial p}{\partial r} + r \frac{\partial^2 p}{\partial r} + \frac{1}{r} \frac{\partial^2 p}{\partial \theta^2} + k^2 r p = 0 \quad (\text{J.18})$$

Utilizando el cambio de variables  $\begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} \frac{\partial p}{\partial r} + \frac{1-2ikr}{2r}p \\ \frac{1}{r} \frac{\partial p}{\partial \theta} \end{bmatrix}$  se obtiene el sistema de ecuaciones diferenciales de primer orden

$$r^2 \frac{\partial u}{\partial r} + r \frac{\partial v}{\partial \theta} + r(ikr + \frac{1}{2})u + (1 - 2ik - 2k^2 r^2)p = 0 \quad (\text{J.19})$$

$$ru - r \frac{\partial p}{\partial r} + \frac{2ikr - 1}{2}p = 0 \quad (\text{J.20})$$

$$rv - \frac{\partial p}{\partial \theta} = 0 \quad (\text{J.21})$$

Vamos a agregar una nueva ecuación diferencial de primer orden combinando las ecuaciones (J.20) y (J.21). Para obtener dicha ecuación de compatibilidad derivamos (J.20) respecto a  $\theta$  y (J.21) respecto a  $r$ , a la ecuación resultante de esta última le restamos la primera multiplicada por  $\frac{1}{r}$ , luego el resultado se multiplica por  $r$  obteniéndose

$$rv + r^2 \frac{\partial v}{\partial r} - r \frac{\partial u}{\partial \theta} - \frac{(2ikr - 1)}{2} \frac{\partial p}{\partial \theta} = 0 \quad (\text{J.22})$$

De igual manera como se desarrolló en la sección anterior se planteará un sistema de ecuaciones que es equivalente al sistema (J.19-J.22). Luego se demostrará que son equivalentes y después se demostrará la elipticidad del último sistema.

Sea el sistema

$$r^2 \frac{\partial u}{\partial r} + r \frac{\partial v}{\partial \theta} + r(ikr + \frac{1}{2})u + (1 - 2ikr - 2k^2 r^2)p = 0 \quad (\text{J.23})$$

$$ru + r \frac{\partial \vartheta}{\partial \theta} - r \frac{\partial p}{\partial r} + \frac{2ikr - 1}{2}p = 0 \quad (\text{J.24})$$

$$rv - r^2 \frac{\partial \vartheta}{\partial r} - \frac{\partial p}{\partial \theta} = 0 \quad (\text{J.25})$$

$$rv + r^2 \frac{\partial v}{\partial r} - r \frac{\partial u}{\partial \theta} - r^2 \frac{\partial \vartheta}{\partial r} - \frac{(2ikr - 1)}{2} \frac{\partial p}{\partial \theta} = 0 \quad (\text{J.26})$$

Para demostrar que el sistema (J.23-J.26) es equivalente al (J.19-J.22) basta que  $\vartheta = \text{Constante}$ , procederemos a probar esto

Primero, derivamos las ecuaciones (J.24) respecto a  $\theta$  y (J.25) respecto a  $r$ . Luego a esta última le restamos  $\frac{1}{r}$  veces la anterior, obteniéndose

$$v + r \frac{\partial v}{\partial r} - \frac{\partial u}{\partial \theta} - \frac{2ikr - 1}{2r} \frac{\partial p}{\partial \theta} - r \frac{\partial \vartheta}{\partial r} - r^2 \frac{\partial^2 \vartheta}{\partial r^2} - r \frac{\partial \vartheta}{\partial r} - \frac{\partial^2 \vartheta}{\partial \theta^2} = 0 \quad (\text{J.27})$$

Utilizando (J.26) obtenemos

$$-r^2 \frac{\partial^2 \vartheta}{\partial r^2} - r \frac{\partial \vartheta}{\partial r} - \frac{\partial^2 \vartheta}{\partial \theta^2} = 0 \quad (\text{J.28})$$

o lo que es lo mismo  $\Delta \vartheta = 0$ . Si se utiliza la condición de frontera  $\frac{\partial \vartheta}{\partial \eta} = 0$  entonces se concluye que  $\vartheta = \text{Constante}$ . Por ello los sistemas (J.19-J.22) y (J.23-J.26) son equivalentes.

Ahora comprobaremos que el sistema (J.23)-(J.26) es elíptico y así (J.19)-(J.22) también lo es, para ello lo escribimos de la forma  $Au = (A_2 \frac{\partial}{\partial \theta} + A_1 \frac{\partial}{\partial r} + A_0)u$  donde

$$A_2 = \begin{bmatrix} 0 & 0 & r & 0 \\ 0 & 0 & 0 & r \\ -1 & 0 & 0 & 0 \\ \frac{1-2ikr}{2} & -r & 0 & 0 \end{bmatrix} \text{ y } A_1 = \begin{bmatrix} 0 & r^2 & 0 & 0 \\ -r & 0 & 0 & 0 \\ 0 & 0 & 0 & -r^2 \\ 0 & 0 & r^2 & -r^2 \end{bmatrix}$$

Para demostrar la elipticidad de un sistema de ecuaciones diferenciales de primer orden expresado en la forma  $Au = (A_2 \frac{\partial}{\partial \theta} + A_1 \frac{\partial}{\partial r} + A_0)u$  hay que probar que

$$\det(A_2 \eta + A_1 \varepsilon) \neq 0 \quad \forall \begin{bmatrix} \varepsilon \\ \eta \end{bmatrix} \neq \begin{bmatrix} 0 \\ 0 \end{bmatrix}$$

En nuestro caso tenemos que  $\det(A_2 \eta + A_1 \varepsilon) = r^3(\eta^2 + r^2 \varepsilon^2)^2 \neq 0 \quad \forall \begin{bmatrix} \varepsilon \\ \eta \end{bmatrix} \neq \begin{bmatrix} 0 \\ 0 \end{bmatrix}$   
con  $r \neq 0$

Así el sistema (J.23)-(J.26) es elíptico si  $r \neq 0$

# Bibliografía

- [Ant] H. Anton and Ch. Rorres, *Elementary Linear Algebra, Applications Version*, 8th ed., Wiley, USA ,p.p. 822, (2000).
- [Azi] A. K. Aziz and A. Werschulz, *On the Numerical Solution of the Helmholtz Equation by the Finite Element Method*, SIAM J. Numer. Anal., Vol 17, **5**, 681-686 (1980).
- [Ba1] I. Babuška, F. Ihlenburg, T. Strouboulis and S. K. Gangaraj, *A Posteriori Error Estimation for Finite Element Solutions of Helmholtz Equation. Part I: The Quality of Local Indicators and Estimators* , International Journal for Numerical Methods in Engineering, **40**, 3443-3462 (1997).
- [Ba2] I. Babuška, F. Ihlenburg, T. Strouboulis and S. K. Gangaraj. *A Posteriori Error Estimation for Finite Element Solutions of Helmholtz Equation. Part II: Estimation of the Pollution Error* , International Journal for Numerical Methods in Engineering, Vol 40, **4**, 789-837 (1998).
- [Bia] B. Bialecki and G. Fairweather. *Orthogonal Spline Collocation Methods for Partial Differential Equations*. Journal of Computational and Applied mathematics. pp. 55-82 (2001).

- [Boc] P.B. Bochev and M. D. Gunzburger, *Finite Element Methods of Least-Squares Type*, SIAM Rev., **40**, 1-7(1990).
- [Bra] J. Bramble, R. Lazarov, and J.E. Pasciak, *Least-Squares for Second Order Elliptic Problems*, Computer Methods in Applied Mechanics and Engineering, **152**, 195-210(1998).
- [Bur] R. Burden and J.D. Faires, *Numerical Analysis* . Brooks/Cole. Six Edition. pp. 811,(1997).
- [Cai] Z. Cai, R. Lazarov, T. A. Manteuffel and S. F. McCormick, *First-Order Least Squares for Second Order Partial Differential Equations: Part I*, SIAM J. Numer. Anal. Vol. 31, **6**, 1785-1799 (1994).
- [Car] G.F. Carey and Y. Shen, *Convergence Studies of Least-Squares Finite Element for First-Order System*, Communications in Applied Numerical Methods, **5**, 427-434(1989).
- [Cha] C.L. Chang, *A Least-Squares Finite Element Method for the Helmholtz Equation*, Computer Methods in Applied Mechanics and Engineering, **83**, 1-7(1990).
- [Coc] B. Cockburn. *Devising Discontinuous Galerkin Methods for Non-linear Hyperbolic Conservation Laws*. Journal of Computational and Applied mathematics. pp. 187-204 (2001).
- [Dah] W. Dahmen. *Wavelet Methods for PDE's - Some Recent Developments*. Journal of Computational and Applied mathematics. pp. 133-186 (2001).
- [Der] A. Deraemaeker, I. Babuška and Ph. Bouillard, *Dispersion and Pollution of the FEM Solution for the Helmholtz Equation in One, Two and Three*

- Dimensions*, International Journal for Numerical Methods in Engineering, Vol. 46, 471-499(1999).
- [Eva] G. Evans, J. Blackledge and P. Yardley *Numerical methods for Partial Differential Equations*. Springer Ungraduate Mathematics Series. pp. 290 (2000).
- [Fla] J. Flaherty. *CSCI, MATH 6860. Finite Element Analysis*. Lecture Notes: Spring. (2000).
- [Fox] Ch. Fox. *An Introduction to Calculus of Variations*. Dover. New York pp. 271 .(1987).
- [Gel] I. M. Gelfand and S. V. Fomin. *Calculus of Variations*. Dover. New York pp. 232 .(2000).
- [Got] D. Gottlieb and J. S. Hesthaven. *Spectral Methods for Hyperbolic Problems*. Journal of Computational and Applied mathematics. pp. 83-132 (2001).
- [Gup] K. Gupta, *A Brief History of the Beginning of Finite Element Method*, International Journal for Numerical Methods in Engineering, **39**, 3761-3774 (1996).
- [Guz] M. de Guzmán, *Ecuaciones Diferenciales Ordinarias, Teoría de Estabilidad y Control*, Alhambra. España. pp. 300 (1975).
- [Ih1] F. Ihlenburg. *Finite Element Analysis of Acoustic Scattering*. Applied Mathematical Sciences, **132**. Springer-Verlag. New York. pp. 224. (1998)
- [Ih2] F. Ihlenburg and I. Babuška. *Dispersion Analysis and Error Estimation of Galerkin Finite Element Methods for the Helmholtz Equation*. International Journal for Numerical Methods in Engineering. **38**, 3745-3774. (1995)

- [Jia] B.N. Jiang, *The Least-Squares Finite Element Method: Theory and Applications in Computational Fluid Dynamics and Electromagnetics*, Scientific Computation, New York, ,p.p. 418, 1998.
- [Joh] C. Johson, *Advanced Methods in Scientific Computing*, Computer Science 522, pp. 165. (1997).
- [Kel] J.B. Keller and D. Givoli, *Exact Non-reflecting Boundary Conditions*, Journal of Computational Physics, **82**, 172-192(1989).
- [Kis] A. Kiseliiov, M. Krasnov y G. Makarenko., *Problemas de Ecuaciones Diferenciales Ordinarias*, Editorial MIR, cuarta reimpresión. URSS. pp. 253. (1988).
- [Le] B. Lee, T. A. Manteuffel, S.F. McCormick, and J. Ruge, *First Order System Least-Squares for the Helmholtz Equation*, SIAM J. Sci. Comput., Vol 21, **5**, 1927-1949 (2000).
- [Lev] H. Levy and F. Lessman. *Finite Difference Equations*. Dover. New York, pp. 278. (1992).
- [Ob] A. Oberai and P. Pinsky, *A Numerical Comparison of the Finite Element Methods for the Helmholtz Equation*, Journal of Computational Acoustic, **1**, 211-221 (2000).
- [Peh] A. I. Pehlivanoh, G. F. Carey and R. D. Lazarov, *Least-Squares Mixed Finite Elements for Second-Order Elliptic Problems*, SIAM J. Numer. Anal. Vol. 31, **5**, 1368-1377 (1994).

- [Ran] R. Rannacher. *Adaptive Galerkin Finite Element Methods for Partial Differential Equations*. Journal of Computational and Applied mathematics. pp. 205-234 (2001).
- [Red] J.N. Reddy, *An Introduction to the Finite Element Method*, 2nd ed., McGraw Hill, Boston, (1993).
- [Str] J. Strikwerda, *Finite Difference Schemes and Partial Differential Equations*, The Wadsworth & Brooks/Cole Mathematics Series. California, USA. pp. 386, (1989).
- [Sul] S. Suleau and Ph. Bouillard, *One Dimensional Dispersion Analysis for the Element-free Galerkin Method for the Helmholtz Equation*, Soviet Math. Dokl. **47**, 1169-1188, (2000).
- [Sur] M. Suri. *The  $p$  and  $hp$  Finite Element Method for Problems on thin Domains*. Journal of Computational and Applied mathematics. pp. 235-260 (2001).
- [Tho] V. Thomée. *From Finite Differences to Finite Elements. A Short History of Numerical Analysis of Partial Differential Equations*. Journal of Computational and Applied mathematics. pp. 1-54 (2001).
- [Wei] E. W. Weisstein. *CRC Concise Encyclopedia of Mathematics*. Chapman & Hall/ CRC. pp. 1969. (1999).