

UNA INTRODUCCIÓN AL USO DE LAS FUNCIONES DE BASE RADIAL PARA
RESOLVER NUMÉRICAMENTE ECUACIONES DIFERENCIALES PARCIALES

DANNA KATHERINE CHÁVEZ SÁNCHEZ

UNIVERSIDAD INDUSTRIAL DE SANTANDER
FACULTAD DE CIENCIAS
ESCUELA DE MATEMÁTICAS
BUCARAMANGA

2021

UNA INTRODUCCIÓN AL USO DE LAS FUNCIONES DE BASE RADIAL PARA
RESOLVER NUMÉRICAMENTE ECUACIONES DIFERENCIALES PARCIALES

DANNA KATHERINE CHÁVEZ SÁNCHEZ

TRABAJO DE GRADO PARA OPTAR AL TÍTULO DE
MATEMÁTICA

DIRECTOR

ÉLDER JESÚS VILLAMIZAR ROA

Ph.D EN MATEMÁTICAS

CODIRECTOR

IVÁN JAVIER SÁNCHEZ GALVIS

Ms. EN INGENIERÍA ELECTRÓNICA

UNIVERSIDAD INDUSTRIAL DE SANTANDER

FACULTAD DE CIENCIAS

ESCUELA DE MATEMÁTICAS

BUCARAMANGA

2021

Dedicado a mi madre Blanca Azucena Sánchez Ortiz, mis hermanos Sergio Andrés Navas Sánchez, Juliana Andrea Navas Sánchez, Daniela Chávez Sánchez, Anyha Isabella Santamaria Sánchez, mis once perros Lucas, Ringo, Mora, Coco, Dobbie, Terry, Lulú, Maxi, Niña, Almendra, Canela, mis cinco gatos, Manchas, Caín, Guasón, Boris, Lilith y mi paloma Lúpe, que son el motor de mi vida.

AGRADECIMIENTOS

Deseo agradecer de forma especial a mis docentes **Élder Jesús Villamizar Roa, Iván Javier Sánchez Galvis** por la dedicación y apoyo que han brindado a este trabajo, por el respeto a mis sugerencias e ideas y por la dirección y el rigor que ha facilitado a las mismas.

Así mismo, agradezco a los profesores que fueron guía de aprendizaje y de ejemplo profesional, la colaboración de administrativos de la Universidad Industrial de Santander, agradezco también al incentivo recibido por parte de Ecopetrol para la realización de este trabajo. Este trabajo es parte del Acuerdo de Cooperación No. 27, derivado del Convenio Marco No. 5222395 suscrito entre Ecopetrol y la Universidad Industrial de Santander.

Un trabajo de grado siempre es fruto de ideas, proyectos y esfuerzos previos pero también es fruto del reconocimiento y apoyo vital que nos ofrecen las personas que nos estiman, sin el cual no tendríamos la fuerza y energía que nos anima a crecer como personas y como profesionales.

Gracias a mi familia, a mis padres, mis hermanos y mis mascotas, porque son el motor de cada meta que tengo y es gracias a ustedes que hoy en día puedo culminar este capítulo en mi vida.

Gracias a mis amigos, **Rosa Ximena Barajas Rincón, Leidy Paola Santos Canticus, Karen Daniela Arana Romero, María Angélica Oliveros, Sergio Andrés Jiménez** que estuvieron acompañándome en los momentos más difíciles, que siempre fueron un apoyo moral y humano necesario para este proyecto y profesión.

Pero sobre todo, gracias a **Edwin René Araque Alarcón** por su paciencia, comprensión y solidaridad con este proyecto. Sin su apoyo no hubiese sido fácil concluir este trabajo y, por eso, este trabajo es también el suyo.

A todos muchas gracias.

CONTENIDO

	pág.
INTRODUCCIÓN	12
1. INTERPOLACIÓN CON FUNCIONES DE BASE RADIAL	15
1.1. FUNCIONES DE BASE RADIAL	16
1.2. FUNCIONES DEFINIDAS POSITIVAS	21
1.3. FUNCIONES CONDICIONALMENTE DEFINIDAS POSITIVAS	27
1.4. ESPACIOS NATIVOS	32
1.5. ESTIMACIONES DE ERROR	38
1.6. PARÁMETRO DE FORMA	50
2. ESQUEMA DE COLOCACIÓN ASIMÉTRICO	53
2.1. PROBLEMA LINEAL ESTACIONARIO	53
2.2. Problema lineal estacionario	56
2.3. ELEMENTOS FINITOS VS MÉTODO DE COLOCACIÓN ASIMÉTRICO	61
3. MÉTODO DE COLOCACIÓN ASIMÉTRICO APLICADO A LA ECUACIÓN DEL CALOR	65
3.1. GENERALIDADES	65
3.2. ANÁLISIS DE ESTABILIDAD	68
3.3. ECUACIÓN LINEAL DEL CALOR	70
3.4. COMPARACIÓN ENTRE EL MÉTODO ELEMENTOS FINITOS Y MÉTODO DE COLOCACIÓN ASIMÉTRICO	79
BIBLIOGRAFÍA	84
4. Número de condición de una matriz	88

LISTA DE FIGURAS

	pág.
Figura 1. Interpolación con FBR de la función $f(x, y) = \sin(x + y)$, usando la función multicuadrática para un valor de $c = 3$.	20
Figura 2. Interpolación con FBR de la función $f(x, y) = \sin(x + y)$, usando la función multicuadrática para un valor de $c = 3$.	21
Figura 3. Distancia de llenado	39
Figura 4. Distribución nodos 2D.	54
Figura 5. Solución analítica del problema de Poisson (50) $u(x, y) = (1 - x^2)(2y^3 - 3y^2 + 1)$.	56
Figura 6. Figura 1: Representa la solución analítica del problema de Poisson (50). Figura 2: Representa la solución numérica del problema de Poisson (50) usando el método asimétrico. Figura 3: Expone el error puntual entre la solución analítica y solución numérica. Figura 4: Describe el dominio para 100 centros distribuidos uniformemente.	58
Figura 7. Método de aproximación estacionaria: N se fija en $N = 100$. Izquierda: Error máximo frente al parámetro de forma del problema de Poisson (azul). Límite de error de la forma 45 (rojo). Centro: Error cuadrático medio contra el parámetro de forma (verde). Límite de error de la forma 45 (rojo). Derecha: Número de condición del sistema (57) vs el parámetro de forma.	59
Figura 8. Método de aproximación no estacionaria: El parámetro de forma se fija en $c = 1.1$. Izquierda: Error máximo frente al número de centros (azul). Límite de error de la forma 45 (rojo). Centro: Error cuadrático medio contra el número de centros N (verde). Límite de error de la forma 45 (rojo). Derecha: Número de condición del sistema (57) vs número de centros	60

- Figura 9. Interpolación del problema de Poisson variando tanto N como c .
Izquierda: Error máximo frente al número de centros. Centro: Error cuadrático medio contra el número de centros N . Derecha: Parámetro de forma vs número de centros. 61
- Figura 10. Dominio rectangular de la superficie con 1541 nodos. 62
- Figura 11. Solución numérica método elementos finitos. Izquierda: Solución analítica problema de Poisson (50)-(52). Centro: Solución numérica del problema de Poisson (50)-(52). Derecha: Error puntual entre la solución analítica y solución numérica del problema (50)-(52). 62
- Figura 12. Solución numérica método de colocación asimétrico con parámetro de forma $c = 0.1$. La imagen se divide en solución analítica (izquierda), solución numérica (centro) del problema de Poisson descrito en (50) y error puntual entre las soluciones antes mencionadas (derecha). 63
- Figura 13. Método de aproximación estacionaria: El número de nodos se fija en $N = 100$, $\Delta t = 0.01$, y $t_{\text{máx}} = 1$. Izquierda: Error máximo frente al parámetro de forma c . Derecha: Error cuadrático medio vs parámetro de forma c . 73
- Figura 14. Método de aproximación estacionaria: El número de nodos se fija en $N = 100$, $t_{\text{máx}} = 1$ y distintos pasos temporales Δt . Izquierda: Error máximo frente al parámetro de forma c . Derecha: Error cuadrático medio vs parámetro de forma c . 74
- Figura 15. Método de aproximación no estacionaria: El parámetro de forma se fija en $c = 0.1$, $\Delta t = 0.01$ y $t_{\text{máx}} = 1$. Izquierda: Error máximo frente al número de centros N . Derecha: Error cuadrático medio vs número de centros N . 74
- Figura 16. Método de aproximación no estacionaria: El parámetro de forma se fija en $c = 0.1$, $t_{\text{máx}} = 1$ y distintos pasos temporales Δt . Izquierda: Error máximo frente al número de centros N . Derecha: Error cuadrático medio vs número de centros N . 75

- Figura 17. Método de aproximación estacionaria: El número de nodos se fija en $N = 100$, $t_{\text{máx}} = 1$ y $\Delta t = 0.01$. Izquierda: Error máximo frente al parámetro de forma c . Derecha: Error cuadrático medio vs parámetro de forma c . 77
- Figura 18. Método de aproximación estacionaria: El número de nodos se fija en $N = 100$, $t_{\text{máx}} = 1$ y distintos pasos temporales Δt . Izquierda: Error máximo frente al parámetro de forma c . Derecha: Error cuadrático medio vs parámetro de forma c . 78
- Figura 19. Método de aproximación no estacionaria: El parámetro de forma se fija en $c = 0.1$, $t_{\text{máx}} = 1$ y distintos pasos temporales Δt . Izquierda: Error máximo frente al número de nodos N . Derecha: Error cuadrático medio vs número de nodos N . 79
- Figura 20. Método de aproximación no estacionaria: El parámetro de forma se fija en $c = 0.5$, $t_{\text{máx}} = 1$ y distintos pasos temporales Δt . Izquierda: Error máximo frente al número de nodos N . Derecha: Error cuadrático medio vs número de nodos N . 79
- Figura 21. Dominio de la superficie discretizado en 1541 nodos. 80
- Figura 22. Solución numérica del problema a través del método de elementos finitos. Izquierda: Solución analítica ecuación del calor (75). Centro: Solución numérica de ecuación de calor (75). Derecha: Error puntual entre la solución analítica y solución numérica del problema (75). 81
- Figura 23. Solución numérica método de colocación asimétrico. Izquierda: Solución analítica ecuación de calor (75). Centro: Solución numérica de ecuación de calor (75). Derecha: Error puntual entre la solución analítica y solución numérica del problema (75). 82

LISTA DE TABLAS

	pág.
Tabla 1. Funciones de base radial comúnmente usadas y su respectivo grado del polinomio.	31
Tabla 2. Estimaciones de error para varios interpolantes FBR en términos de distancia de llenado h .	48
Tabla 3. Comparación método de colocación asimétrico vs método elementos finitos.	63
Tabla 4. Comparación método colocación asimétrico vs método elementos finitos.	63
Tabla 5. Variación del parámetro c para el esquema implícito con $N = 100$, $\Delta t = 0.01$, $t_{\text{máx}} = 1$.	73
Tabla 6. Variación del número de nodos N para el esquema implícito con $c = 0.1$, $\Delta t = 0.01$, $t_{\text{máx}} = 1$.	75
Tabla 7. Variación del parámetro c para el esquema explícito con $N = 100$, $\Delta t = 0.01$, $t_{\text{máx}} = 1$.	77
Tabla 8. Comparación entre la variación del tamaño temporal en el método elementos finitos, y método de colocación asimétrico (esquema implícito), la disminución de los errores ECM , $E_{\text{máx}}$ y el tiempo de cómputo de la implementación para ambos métodos numéricos.	83

RESUMEN

TÍTULO: UNA INTRODUCCIÓN AL USO DE LAS FUNCIONES DE BASE RADIAL PARA RESOLVER NUMÉRICAMENTE ECUACIONES DIFERENCIALES PARCIALES. *

AUTOR: DANNA KATHERINE CHÁVEZ SÁNCHEZ **

PALABRAS CLAVE: MÉTODOS LIBRES DE MALLA, INTERPOLACIÓN CON FUNCIONES DE BASE RADIAL, MÉTODO DE COLOCACIÓN ASIMÉTRICO, CONVERGENCIA EXPONENCIAL, FUNCIÓN RADIAL MULTICUADRÁTICA, MÉTODO DE ELEMENTOS FINITOS.

DESCRIPCIÓN:

Este trabajo está relacionado con el análisis numérico de soluciones de ecuaciones diferenciales parciales mediante funciones de base radial. Este trabajo consta de tres partes fundamentales; el primero corresponde a una revisión de definiciones y resultados preliminares sobre la interpolación con funciones de base radial, las funciones positivas definidas y las funciones definidas condicionalmente, así como algunas de sus caracterizaciones. A partir de la generalización de las funciones mencionadas, se estudia el espacio funcional asociado, denominado « espacio nativo », que se requiere para establecer la convergencia del interpolador. La segunda parte está dedicada a la descripción del método de colocación asimétrica para un problema lineal general. Este método fue propuesto por E. J. Kansa en 1990¹², el cual ha sido utilizado en la aproximación de soluciones para ecuaciones diferenciales parciales mediante funciones de base radial. Además, el orden de convergencia exponencial del método de colocación asimétrico se verifica utilizando la función radial multicuadrática en un problema de Poisson bidimensional con condiciones de contorno mixtas. Finalmente, en la tercera parte se analiza el Método de Colocación Asimétrico para la ecuación de calor lineal en el escenario bidimensional, basado en la función radial multicuadrática, y con un esquema implícito y explícito en la discretización temporal. Además, en ambas partes de la implementación numérica utilizando el Método de Colocación Asimétrico, se realiza una comparación entre el método sin malla y el método clásico con elementos finitos.

* Trabajo de grado

** Facultad de Ciencias. Escuela de Matemáticas. Director: Élder Jesús Villamizar Roa, Ph.D en Matemáticas. Codirector: Iván Javier Sánchez Galvis, Ms en ingeniería electrónica.

¹ E. J. Kansa. "Multiquadrics—a scattered data approximation scheme with applications to computational fluid-dynamics. I. surface approximations and partial derivative estimates," en: *Computers and Mathematics with Applications*, 19(8-9) (1990), págs. 127-145.

² E. Kansa. "Multiquadrics—a scattered data approximation scheme with applications to computational fluid-dynamics. II. solutions to parabolic, hyperbolic and elliptic partial differential equations," en: *Computers and Mathematics with Applications*, 19(8-9) (1990), 147–161.

ABSTRACT

TITLE: AN INTRODUCTION TO THE USE OF RADIAL BASIS FUNCTIONS FOR THE NUMERICAL SOLUTION OF PARTIAL DIFFERENTIAL EQUATIONS. *

AUTHOR: DANNA KATHERINE CHÁVEZ SÁNCHEZ **

KEYWORDS: MESH-FREE METHOD, INTERPOLATION WITH RADIAL BASIS FUNCTIONS, ASYMMETRIC COLLOCATION METHOD, EXPONENTIAL CONVERGENCE, MULTIQUADRATIC RADIAL FUNCTION, FINITE ELEMENTS METHOD.

DESCRIPTION:

This work is related to the numerical analysis of solutions of partial differential equations through radial basis functions. This work consists of three fundamental parts; the first one corresponds to a review of definitions and preliminaries results on the interpolation with radial basis functions, the defined positive functions and conditionally defined functions, as well as some of their characterizations. From the generalization of the mentioned functions, we study the associated function space, denominated «Native space», which is required to establish the convergence of the interpolator. The second part is devoted to the description of the Asymmetric Collocation Method for a general linear problem. This method was proposed by E. J. Kansa in 1990¹², which has been used in the approximation of solutions for partial differential equations by means of radial basis functions. Moreover, the order of exponential convergence of the asymmetric collocation method is verified using the multi-quadratic radial function in a two-dimensional Poisson problem with mixed boundary conditions. Finally, in the third part, the Asymmetric Collocation Method is analyzed for the linear heat equation in the bi-dimensional setting, based on the multi-quadratic radial function, and with an implicit and explicit scheme in the temporal discretization. Besides, in both parts of the numerical implementation using the Asymmetric Collocation Method, a comparison between the mesh-free method and the classical method finite elements is carried out.

* Bachelor Thesis

** Facultad de Ciencias. Escuela de Matemáticas. Director: Élder Jesús Villamizar Roa, Ph.D en Matemáticas. Codirector: Iván Javier Sánchez Galvis, Ms en ingeniería electrónica.

INTRODUCCIÓN

La presente monografía tiene por objetivo hacer una introducción al método de las funciones de base radial (FBR) para la resolución numérica de ecuaciones en derivadas parciales. Particularmente, se hará énfasis en el llamado método asimétrico de Kansas para resolver ecuaciones elípticas y parabólicas.

Tradicionalmente, la solución numérica de ecuaciones diferenciales parciales se ha llevado a cabo a través de métodos de diferencias finitas, elementos finitos y volúmenes finitos, que se basan en la discretización del dominio mediante el uso de mallado. Sin embargo, en el tratamiento de problemas reales, como por ejemplo problemas del estudio de suelos en geología, surgen mallados con geometrías de alta complejidad que requieren gran esfuerzo computacional, generando desventajas en aplicaciones de simulación. Es por ello que en las últimas dos décadas ha cobrado relevancia la investigación y uso de métodos numéricos libres de malla, que como su nombre lo indica, no requieren de una malla para obtener soluciones numéricas de una ecuación diferencial parcial.

En términos generales, la idea de los métodos libres de malla es construir la aproximación del problema a ser analizado, a través de un conjunto de nodos en donde no se necesite una relación entre ellos; es decir, no requieren de una malla. Este procedimiento es ventajoso a la hora de capturar adecuadamente deformaciones en los dominios sobre los cuales se define una ecuación diferencial parcial (EDP). La motivación del uso de métodos libres de malla es reducir los tiempos de cómputo y que el orden de convergencia sea mayor o igual al de los métodos numéricos tradicionales, entendiendo por convergencia al decrecimiento del error obtenido conforme N crece, siendo N el número de nodos. Al día de hoy se conoce que en problemas de EDP con dominios que presentan largas deformaciones y distorsiones, los métodos libres de mallas son bastante apropiados ¹.

¹ A. I. Tolstykh y D. A. Shirobokov. "On using radial basis functions in a finite difference mode with

Dentro de los métodos libres de malla se han desarrollado esquemas numéricos basados en funciones de base radial. Las funciones de base radial originalmente han sido utilizadas como una herramienta efectiva para la interpolación de datos no equiespaciados en $\mathbb{R}^d$, y poseen la característica de que no se requiere de una malla para construir el interpolante, y el problema de interpolación no crece en función de la dimensión espacial en donde están definidos los datos. Adicionalmente, para cierto tipo de problemas, la interpolación con funciones de base radial numéricamente ha mostrado tener una tasa de convergencia superior a la de los métodos tradicionales ². Dentro de los esquemas numéricos basados en funciones de base radial se destacan el método de colocación asimétrico ¹², el método de colocación simétrico ³, el método de soluciones fundamentales ⁴, esquemas de tipo Galerkin ⁵, entre otros.

En este trabajo hacemos un análisis disertativo de la fundamentación teórica de los métodos de aproximación numérica basados en funciones de base radial, y sus aplicaciones para resolver ecuaciones en derivadas parciales. Por la extensión del tema, centramos nuestro análisis en el método de colocación asimétrico, y su aplicación en el análisis de la ecuación de Poisson y la ecuación del calor. Las implementaciones numéricas mediante el método de colocación asimétrico son realizadas usando el software “Matlab”.

applications to elasticity problems,” en: *Computational Mechanics*, 33(1) (2003), págs. 68-79.

- ² J. Muñoz. *Solución numérica de ecuaciones en derivadas parciales mediante funciones radiales (tesis doctoral)*,
- ³ G. E. Fasshauer. *Boundary Elements XXVII: Incorporating Electrical Engineering and Electromagnetics, chapter RBF Collocation Methods and Pseudospectral Methods*,
- ⁴ M.A. Golberg y C.S. Chen. *Boundary Integral Methods - Numerical and Mathematical Aspects, chapter The method of fundamental solutions for potential, Helmholtz and diffusion problems, pages 103–176*,
- ⁵ H. Wendland. “Meshless Galerkin methods using radial basis functions,” en: *Mathematics of Computation*, 68(228) (1999), págs. 1521-1531.

El desarrollo de esta monografía consta de 3 capítulos. En el Capítulo 1 se exponen fundamentos teóricos para la interpolación con funciones de base radial; se aborda un análisis de las funciones definidas positivas, funciones condicionalmente definidas positivas, espacios nativos, entre otros conceptos, que soportan el análisis numérico del método. En el Capítulo 2 hacemos una descripción del método de colocación asimétrico para problemas elípticos lineales estacionarios. Analizamos numéricamente la relación entre el parámetro de forma de la función multicuadrática y el número de nodos. En el Capítulo 3 se describe el método de colocación asimétrico para problemas evolutivos, más específicamente, para dar solución a la ecuación lineal del calor; finalmente, se presenta una comparación entre el método elementos finitos y método de colocación asimétrico para el problema evolutivo, mostrando algunas simulaciones de Matlab.

1. INTERPOLACIÓN CON FUNCIONES DE BASE RADIAL

En diferentes problemas prácticos nos enfrentamos al problema de reconstruir una cierta función desconocida a través de un conjunto de datos. Estos datos corresponden a un conjunto finito de mediciones, y se busca obtener información del proceso estudiado en puntos diferentes a los dados, es decir, se busca una función que sea un buen ajuste de los datos dados y del proceso físico asociado a ellos. En términos más precisos, se busca encontrar una función continua s que interpole los datos, es decir, se quiere hallar $s : \mathbb{R}^d \rightarrow \mathbb{R}$ continua tal que

$$s(x_i) = f_i, \quad i = 1, 2, \dots, N, \quad (1)$$

donde son $(x_i, f_i) \in \mathbb{R}^d \times \mathbb{R}$ son puntos conocidos. Recordemos que en el caso unidimensional, para puntos dados $x_1 < x_2 < \dots < x_N$ y valores $f_1, \dots, f_N$ de una función desconocida f , existe exactamente un polinomio $s \in \mathbb{P}_{N-1}(\mathbb{R})$ que interpola los puntos (x_i, f_i) , $i = 1, \dots, N$. En este caso el espacio de funciones interpolantes es el espacio de polinomios de grado a lo más $N - 1$ y es de resaltar que el espacio de funciones interpolantes no depende ni de los puntos x_i ni de los valores f_i , sino solo del número de puntos. Entonces, es natural pensar en espacios con esta propiedad en el caso d -dimensional; desafortunadamente, si se trabaja en dimensiones mayores que 1, no es posible encontrar un espacio fijo N -dimensional que cumpla con la propiedad de interpolación previamente descrita, hecho conocido en la literatura como Teorema de Mairhuber-Curtis (ver Teorema 1.4).

Una manera conveniente de superar esta dificultad es considerar interpolación con funciones radiales cuya naturaleza no toma en cuenta la dimensión del espacio donde se encuentran los datos. De manera más explícita, se considera que la función s es

una combinación lineal de la forma

$$s(x) = \sum_{j=1}^N \lambda_j \Phi(x - x_j), \quad x \in \mathbb{R}^N, \quad (2)$$

donde $\lambda = [\lambda_1, \lambda_2, \dots, \lambda_N]^T$ es el vector de incógnitas, $\Phi : \mathbb{R}^d \rightarrow \mathbb{R}$ corresponde a la función radial (ver Definición 1.1), $x_j \in \mathbb{R}^N$ son llamados centros de la función. El sistema matricial resultante de (2) al evaluar en los centros de la función y usar (1) queda expresado de la siguiente manera:

$$\begin{bmatrix} \Phi(x_1 - x_1) & \Phi(x_1 - x_2) & \cdots & \Phi(x_1 - x_N) \\ \Phi(x_2 - x_1) & \Phi(x_2 - x_2) & \cdots & \Phi(x_2 - x_N) \\ \vdots & \vdots & \ddots & \vdots \\ \Phi(x_N - x_1) & \Phi(x_N - x_2) & \cdots & \Phi(x_N - x_N) \end{bmatrix} \begin{bmatrix} \lambda_1 \\ \lambda_2 \\ \vdots \\ \lambda_N \end{bmatrix} = \begin{bmatrix} f_1 \\ f_2 \\ \vdots \\ f_N \end{bmatrix},$$

es decir, tenemos un sistema lineal

$$A\lambda = f, \quad (3)$$

donde $A = \{\Phi(x_i - x_j)\}_{1 \leq i, j \leq N} \in \mathbb{R}^{N \times N}$, $f = [f_1, \dots, f_N]^T$, y $\lambda = [\lambda_1, \dots, \lambda_N]^T \in \mathbb{R}^N$.

Para determinar la solución del sistema lineal de ecuaciones (3) requerimos que la matriz A sea no singular, lo que implica determinar qué tipo de funciones radiales hacen que el sistema (3) sea no singular.

En este capítulo mostraremos cómo funciona la interpolación con funciones de base radial y daremos algunos resultados que explican propiedades importantes relativas a la no singularidad de la matriz de interpolación A .

1.1. FUNCIONES DE BASE RADIAL

A continuación presentamos el concepto de funciones de base radial, y definiremos los espacios de Haar, con el fin de poder visualizar el Teorema de Mairhuber-Curtis, que nos brinda condiciones sobre los datos para que sea posible una interpolación

con polinomios definidos sobre $\mathbb{R}^d$.

Definición 1.1. Una función $\Phi : \mathbb{R}^d \rightarrow \mathbb{R}$ es llamada radial, si existe una función univariada $\phi : [0, \infty) \rightarrow \mathbb{R}$ tal que

$$\Phi(x) = \phi(r), \quad (4)$$

donde $r = \|x\|$ y $\|\cdot\|$ es la norma euclidiana en $\mathbb{R}^d$.

A partir de la definición de función radial se puede observar que éstas dependen únicamente de la distancia euclídea entre el argumento y el origen, lo que muestra una propiedad de simetría radial de la función. En consecuencia, la interpolación con dichas funciones no toma en cuenta la dimensión del espacio en donde se encuentran los datos. A continuación se presentan algunos ejemplos de funciones de base radial que son de uso frecuente en aplicaciones.

- Función de distancia

$$\phi(r) = r. \quad (5)$$

- Función gaussiana

$$\phi(r) = e^{-\epsilon r^2}, \quad \epsilon > 0. \quad (6)$$

- Función spline poliarmónico

$$\phi(r) = r^{2k-1}, \quad k \in \mathbb{N}, \quad (7)$$

$$\phi(r) = r^{2k} \ln(r), \quad k \in \mathbb{N}. \quad (8)$$

- Función multicuadrática

$$\phi(r) = \sqrt{c^2 + r^2}, \quad c > 0. \quad (9)$$

- Función multicuadrática inversa

$$\phi(r) = \frac{1}{\sqrt{c^2 + r^2}}, \quad c > 0. \quad (10)$$

Para las funciones multicuadrática y multicuadrática inversa, la constante c es en principio arbitraria.

Como ya se mencionó, un aspecto importante de la interpolación es el de encontrar una función continua s , tal que sea una buena aproximación de un conjunto de datos (x_i, f_i) , $i = 1, \dots, N$, donde el dominio de la función interpolante sea un subconjunto de un espacio de dimensión mayor que 1. Los espacios vectoriales sobre los cuales las funciones interpolantes dependen sólo del número de puntos mas no de los valores explícitos (x_i, f_i) , se denominan espacios de Haar.

Definición 1.2. *Suponga que $\Omega \subseteq \mathbb{R}^d$ contiene al menos N puntos. Sea $V \subseteq C(\Omega)$ un espacio lineal N -dimensional. Entonces V es llamado un espacio de Haar de dimensión N en Ω , si para puntos distintos arbitrarios $x_1, x_2, \dots, x_N \in \Omega$ y $f_1, f_2, \dots, f_N \in \mathbb{R}$ arbitrarios, existe exactamente una función $s \in V$ con $s(x_i) = f_i$, $1 \leq i \leq N$.*

Se sigue de la Definición 1.2 que conjuntos de polinomios de grado $N - 1$ definidos en una variable forman un espacio de Haar N -dimensional. Los espacios de Haar, pueden ser caracterizados de varias maneras, como se muestra a continuación.

Proposición 1.3. *Bajo las condiciones de la Definición 1.2 las siguientes proposiciones son equivalentes:*

1. *V es un espacio de Haar de dimensión N .*
2. *Todo $u \in V \setminus \{0\}$ tiene a lo más $N - 1$ ceros.*
3. *Para cualquier conjunto de puntos distintos $\{x_1, \dots, x_N\} \subseteq \Omega$ y cualquier base $\{u_1, \dots, u_N\}$ de V , se tiene que*

$$\det A \neq 0, \text{ con } A = (u_j(x_k)), \quad j, k = 1, \dots, N.$$

Demostración. Sean V un espacio de Haar de dimensión N y $u \in V \setminus \{0\}$. Suponga que u tiene N ceros, digamos $x_1, \dots, x_N$. En este caso, la función u y la función nula interpolan los N puntos. Por unicidad, se sigue que $u \equiv 0$, lo que contradice la suposición. Ahora veamos que 2. implica 3. Si $\det A = 0$, existe un vector $\alpha \in \mathbb{R}^N \setminus \{0\}$, con

$A\alpha = 0$, es decir,

$$\sum_{j=1}^N \alpha_j u_j(x_k) = 0, \text{ para } 1 \leq k \leq N, \alpha = (\alpha_1, \dots, \alpha_N).$$

Esto significa que la función $u = \sum_{j=1}^N \alpha_j u_j$ tiene N ceros, y por hipótesis, u debe ser idénticamente cero. Sin embargo, esto es imposible dado que $\alpha \neq 0$ y los u_j forman una base. Finalmente, veamos que 3. implica 1. Sea $A = (u_j(x_k))$ $\det A \neq 0$, podemos considerar la función $u = \sum_{j=1}^N \alpha_j u_j$ como interpolante. La condición de interpolación implica que $\sum_{j=1}^N \alpha_j u_j(x_k) = f_k$, $1 \leq k \leq N$. Naturalmente, los coeficientes, y así u , son únicamente determinadas debido a que A es no singular. $\square$

Por su parte, Mairhuber (1956) y Curtis (1959) demostraron que no es posible llevar a cabo una interpolación con polinomios de grado N en varias variables para un conjunto arbitrario de puntos en $\mathbb{R}^d$. Este es el contenido del siguiente teorema.

Teorema 1.4. ⁶(Mairhuber-Curtis) *Suponga que $\Omega \subseteq \mathbb{R}^d$, $d \geq 2$, contiene un punto interior. Entonces no existe un espacio de Haar sobre Ω de dimensión $N \geq 2$.*

Demostración. Sea $d \geq 2$ y asuma que V es un espacio de Haar sobre Ω con base $\{u_1, \dots, u_N\}$. Como Ω contiene un punto interior, existe una bola $B(x_0, \delta) \subseteq \Omega$ con radio $\delta > 0$. Podemos fijar un conjunto de puntos distintos $\{x_1, x_2, \dots, x_N\} \subset B(x_0, \delta)$, y denotemos por A a la matriz con entradas $A_{jk} = u_k(x_j)$, $j, k = 1, \dots, N$. Entonces, por la definición de espacio de Haar, se tiene que

$$\det(A) \neq 0. \tag{11}$$

Ahora consideramos un camino cerrado P en Ω conectando solo en x_1 y x_2 . Esto es imposible en $\mathbb{R}$, pero es posible en Ω con $d \geq 2$. Entonces podemos cambiar las posiciones de x_1 y x_2 moviendo continuamente estos puntos a lo largo del camino P (sin interferir con cualquier otro punto x_j). Esto significa, sin embargo, que las filas 1 y 2 del determinante (11) han cambiado, lo que lleva a un cambio de signo en el

⁶ G. E. Fasshauer. *Meshfree Approximation Methods with MATLAB*, World Scientific,

determinante.

Como el determinante es una función continua de x_1 y x_2 , se debe tener que $\det(A) = 0$ en algún punto a lo largo de P , lo que contradice (11). $\square$

Por el Teorema 1.4 sabemos que no hay espacios de Haar en el contexto multivariado. Por lo tanto, una manera de interpolar los puntos (x_i, f_i) , es tomar una función radial fija $\Phi : \mathbb{R}^d \rightarrow \mathbb{R}$ y formar el interpolante como en (2).

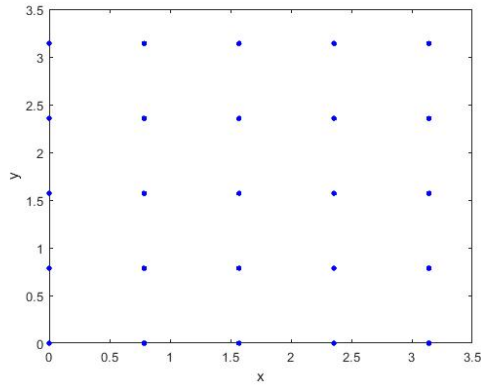
Finalizamos esta sección con el siguiente ejemplo que ilustra el proceso de interpolación por funciones de base radial.

Ejemplo 1.5. En el siguiente ejemplo nos restringiremos al caso bidimensional. Considere la siguiente función

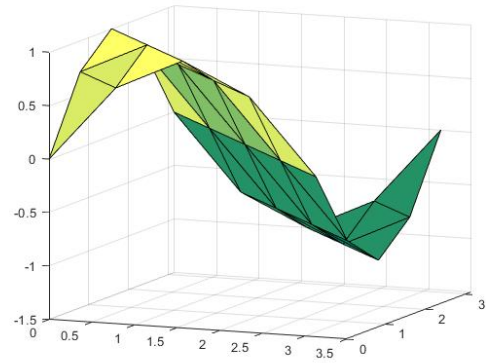
$$f(x, y) = \sin(x + y)$$

definida en el dominio $[0, \pi] \times [0, \pi]$. Tomemos la función básica $\phi(r) = \sqrt{c^2 + r^2}$.

Figura 1. Interpolación con FBR de la función $f(x, y) = \sin(x + y)$, usando la función multicuadrática para un valor de $c = 3$.

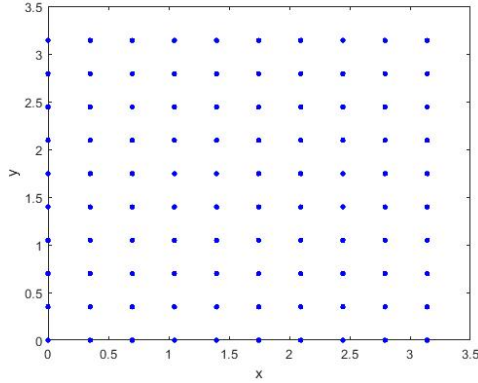


(a) Dominio de la superficie para 25 centros.

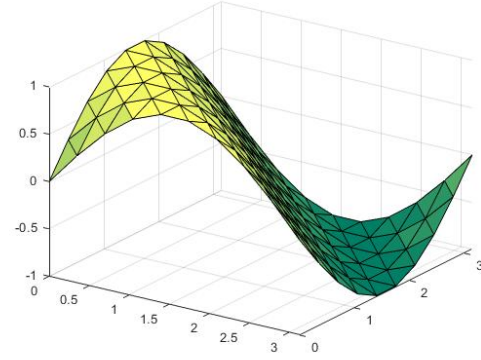


(b) Superficie Interpolada.

Figura 2. Interpolación con FBR de la función $f(x, y) = \sin(x + y)$, usando la función multicuadrática para un valor de $c = 3$.



(a) Dominio de la superficie para 100 centros.



(b) Superficie Interpolada.

Si en el Ejemplo 1.5 consideramos un conjunto de puntos sustancialmente grande, nos veremos enfrentados a la dificultad de computar la inversa de la matriz A en (3). Es por ello que, desde el punto de vista numérico, es conveniente determinar condiciones que garanticen la no singularidad de A .

1.2. FUNCIONES DEFINIDAS POSITIVAS

La siguiente sección describe algunas condiciones que son necesarias para determinar el tipo de función radial que permite la no singularidad de la matriz A del sistema (3).

Definición 1.6. ⁶ Una matriz real simétrica A de orden N es llamada semi definida positiva si su asociada forma cuadrática es no negativa, es decir, si

$$\sum_{j=1}^N \sum_{k=1}^N \alpha_j \alpha_k A_{j,k} \geq 0, \quad (12)$$

para cualquier $\alpha = [\alpha_1, \dots, \alpha_N]^T \in \mathbb{R}^N$. Si el único vector α que satisface la igualdad en (12) es el vector nulo, entonces A es llamada definida positiva.

Una propiedad importante de las matrices definidas positivas es que todos sus valores propios son positivos, y por consiguiente, son matrices invertibles. Por lo tanto, si las

funciones básicas Φ definidas en la Sección 1.1 generan una matriz definida positiva, siempre tendríamos un problema de interpolación bien planteado. A continuación introducimos el concepto de función semi definida positiva y definida positiva.

Definición 1.7. Una función continua $\Phi : \mathbb{R}^d \rightarrow \mathbb{R}$ es llamada semi definida positiva si para cualquier $N \in \mathbb{N}$, todo conjunto de puntos distintos dos a dos $\{x_1, \dots, x_N\} \subset \mathbb{R}^d$, y todo $\alpha = [\alpha_1, \dots, \alpha_N]^T \in \mathbb{R}^N$, verifica que su asociada forma cuadrática es no negativa

$$\sum_{j=1}^N \sum_{k=1}^N \alpha_j \alpha_k \Phi(x_j - x_k) \geq 0. \quad (13)$$

Además, la función Φ es llamada definida positiva en si el único vector α que verifica la igualdad (13) es $\alpha = 0$.

De la Definición 1.7 surge una observación importante, y es que, aunque estamos interesados en problemas con funciones de valor real, esta definición se puede extender para funciones de valores complejos, lo que permite obtener algunas propiedades sobre las funciones semi definidas positivas, incluyendo el famoso Teorema de Bochner (ver Teorema 1.11) que caracteriza las funciones semi definidas positivas en términos de la Transformada de Fourier.

Ejemplo 1.8. Sea $\Phi : \mathbb{R} \rightarrow \mathbb{R}$ definida como

$$\Phi(x) = \begin{cases} 1 & \text{si } x = 0, \\ 0 & \text{caso contrario.} \end{cases}$$

Entonces Φ es definida positiva. En efecto, dado $N \in \mathbb{N}$ y $\{x_1, \dots, x_N\} \subseteq \mathbb{R}$, se tiene que

$$\sum_{j=1}^N \sum_{k=1}^N \alpha_j \alpha_k \Phi(x_j - x_k) = \sum_{j=1}^N |\alpha_j|^2 \geq 0, \text{ para todo escalar } \alpha_j, \alpha_k \in \mathbb{R}.$$

Algunas propiedades que tienen las funciones semi definidas positivas y las funciones definidas positivas se pueden observar a continuación.

Teorema 1.9. ⁶ Suponga que Φ es una función semi definida positiva. Entonces se satisfacen las siguientes propiedades:

1. $\Phi(0) \geq 0$.
2. $\Phi(-x) = \overline{\Phi(x)}$, donde $\overline{\Phi(x)}$ denota el conjugado de $\Phi(x)$.
3. Φ es acotada, más aún, $|\Phi(x)| \leq \Phi(0)$ para todo $x \in \mathbb{R}^d$.
4. $\Phi(0) = 0$ si, y solo si, $\Phi \equiv 0$.
5. Si $\Phi_1, \dots, \Phi_n$ son funciones semi definidas positivas y $b_j \geq 0$, $1 \leq j \leq n$, entonces $\Phi := \sum_{j=1}^n b_j \Phi_j$ es también semi definida positiva. Si existe j_0 , $1 \leq j_0 \leq n$ tal que Φ_{j_0} es definida positiva, y b_{j_0} es positivo, entonces Φ también es definida positiva.
6. El producto de funciones semi definidas positivas (definidas positivas), es semi definida positiva (definida positiva).

Ejemplo 1.10. Considere la función

$$\cos(x) = \frac{1}{2}(e^{ix} + e^{-ix}). \quad (14)$$

Tomemos la función $\Phi : \rightarrow \mathbb{C}$ definida por $\Phi(x) = e^{ix \cdot y}$, con $y \in$ fijo, en donde $x \cdot y$ denota el producto interno usual en $\mathbb{R}^d$. En este caso, dado un conjunto de puntos distintos dos a dos $\{x_1, \dots, x_N\} \subset \mathbb{R}^d$, y $\alpha = [\alpha_1, \dots, \alpha_N]^T \in \mathbb{R}^N$, la forma cuadrática dada definida en (13), verifica que

$$\begin{aligned} \sum_{j=1}^N \sum_{k=1}^N \alpha_j \alpha_k \Phi(x_j - x_k) &= \sum_{j=1}^N \sum_{k=1}^N \alpha_j \alpha_k e^{i(x_j - x_k) \cdot y} \\ &= \sum_{j=1}^N \alpha_j e^{ix_j \cdot y} \sum_{k=1}^N \alpha_k e^{-ix_k \cdot y} \\ &= \left| \sum_{j=1}^N \alpha_j e^{ix_j \cdot y} \right|^2 \geq 0. \end{aligned}$$

De lo anterior se concluye que Φ es semi definida positiva en . Usando la propiedad 5 del Teorema 1.9, $\cos(x)$ es semi definida positiva.

Uno de los resultados más relevantes sobre funciones semi definidas positivas es su caracterización en términos de transformadas de Fourier establecidas por Bochner en 1932 – 1933.

Teorema 1.11. ⁷ (Bochner) *Una función continua $\Phi : \mathbb{R}^d \rightarrow \mathbb{R}$ es semi definida positiva si, y solo si, es la transformada de Fourier de una medida de Borel no negativa finita μ en $\mathbb{R}^d$, es decir,*

$$\Phi(x) = \widehat{\mu}(x) = (2\pi)^{-d/2} \int_{\mathbb{R}^d} e^{ix \cdot \omega} d\mu(\omega), \quad x \in \mathbb{R}^d.$$

Una caracterización completa de las funciones definidas positivas lleva a una discusión muy técnica de las medidas de Borel involucradas que no abarca este documento. Sin embargo, enunciaremos un par de resultados que resultan útiles en la caracterización de este tipo de funciones. El lector interesado en el tema puede tomar como referencia el artículo ⁸.

De la simetría entre la transformada de Fourier y la transformada inversa se puede encontrar un criterio para verificar si una función es definida positiva, el cual se enuncia en el siguiente teorema.

Teorema 1.12. ⁷ *Si $f \in L^1(\mathbb{R}^d)$ es continua, no negativa y no nula, entonces*

$$\Phi(x) := \int_{\mathbb{R}^d} f(w) e^{-ix \cdot w} dw, \quad x \in \mathbb{R}^d,$$

es definida positiva.

Ejemplo 1.13. *Sea $\Phi(x) = e^{-\epsilon \|x\|_2^2}$, $\epsilon > 0$. Entonces Φ es definida positiva en $\mathbb{R}^d$.*

Demostración. En primer lugar mostremos que la función $G(x) := e^{-\|x\|_2^2/2}$ es definida

⁷ Holguer. Wendland. *Scattered Data Approximation, Cambridge Monographs on Applied and Computational Mathematics*,

⁸ K. F. Chang. “Strictly positive definite functions”. En: *J. Approx. Theory* 87 (1996), págs. 148-158.

positiva. Note que $\widehat{G} = G$. En efecto,

$$\begin{aligned}\widehat{G}(x) &= (2\pi)^{-d/2} \int_{\mathbb{R}^d} e^{-\|w\|_2^2/2} e^{-ix \cdot w} dw \\ &= \prod_{j=1}^d \left[(2\pi)^{-1/2} \int_{-\infty}^{\infty} e^{-w_j^2/2} e^{-ix_j w_j} dw_j \right].\end{aligned}\tag{15}$$

Si $g(t) = e^{-t^2/2}$, entonces

$$\begin{aligned}\widehat{g}(r) &= (2\pi)^{-1/2} \int_{-\infty}^{\infty} e^{-(t^2/2 + irt)} dt \\ &= (2\pi)^{-1/2} e^{-r^2/2} \int_{-\infty}^{\infty} e^{-(t+ir)^2/2} dt \\ &= (2\pi)^{-1/2} e^{-r^2/2} \int_{-\infty}^{\infty} e^{-t^2/2} dt \\ &= (2\pi)^{-1/2} e^{-r^2/2} (2\pi)^{1/2} \\ &= e^{-r^2/2}.\end{aligned}\tag{16}$$

De (15) y (16) se tiene que

$$\widehat{G}(x) = e^{-\|x\|_2^2/2} = G(x).$$

Definamos ahora $G_\epsilon(x) = G(x/\epsilon)$. Entonces $G_{1/\sqrt{2\epsilon}}(x) = \Phi(x)$ y $\widehat{G}_\epsilon(x) = (\epsilon)^d \widehat{G}(\epsilon x)$.

Así, obtenemos que

$$\begin{aligned}\Phi(x) &= G_{1/\sqrt{2\epsilon}}(x) = (2\epsilon)^{-d/2} G(x/\sqrt{2\epsilon}) \\ &= (2\epsilon)^{-d/2} (2\pi)^{-d/2} \int_{\mathbb{R}^d} e^{-\|\omega\|_2^2/(4\epsilon)} e^{-x\omega} d\omega \\ &= (2)^{-d} (\epsilon\pi)^{-d/2} \int_{\mathbb{R}^d} e^{-\|\omega\|_2^2/(4\epsilon)} e^{-x\omega} d\omega,\end{aligned}$$

y así, por el Teorema 1.12, Φ es definida positiva. □

El Teorema 1.12, aunque útil, supone que la función f sea integrable en $\mathbb{R}^d$, lo cual puede ser restrictivo en el sentido que deja fuera a funciones de base radial muy importantes, tales como la función multicuadrática $\Phi(x) = \sqrt{c^2 + \|x\|^2}$. En este sentido, es conveniente generalizar el concepto de función definida positiva, dando lugar a las

funciones condicionalmente definidas positivas, definidas en la siguiente sección. Previo a la definición de funciones condicionalmente definidas positivas, enunciamos el concepto de funciones completamente monótonas.

Definición 1.14. *Una función $\phi : [0, \infty) \rightarrow \mathbb{R}$ que pertenece a $C[0, \infty) \cap C^\infty(0, \infty)$, y satisface*

$$(-1)^l \phi^{(l)}(r) \geq 0, \quad r > 0, \quad l = 0, 1, 2, \dots$$

es llamada completamente monótona en $[0, \infty)$.

Ejemplo 1.15. *La función $\phi(r) = e^{-\epsilon r}$, $\epsilon \geq 0$, es completamente monótona en $[0, \infty)$.*

Ejemplo 1.16. *La función $\phi(r) = e^{-\epsilon r}$, $\epsilon \geq 0$, es completamente monótona en $[0, \infty)$. Note que*

$$(-1)^l \phi^{(l)}(r) = \epsilon^l e^{-\epsilon r} \geq 0, \quad l \in \mathbb{N}.$$

Ejemplo 1.17. *La función $\phi(r) = \frac{1}{(1+r)^\beta}$, $\beta \geq 0$, es completamente monótona en $[0, \infty)$ ya que*

$$(-1)^l \phi^{(l)}(r) = (-1)^l \beta(\beta+1) \cdots (\beta+l-1)(1+r)^{-\beta-l} \leq 0, \quad l \in \mathbb{N}.$$

La siguiente es una caracterización de las funciones completamente monótonas en relación con las funciones radiales definidas positivas. Hasta donde sabemos, esta relación fue señalada por primera vez por Schoenberg en 1938, dando así, condiciones suficientes para la invertibilidad de la matriz A en (3). Es decir, funciones que satisfagan las hipótesis del siguiente teorema, pueden ser utilizadas para el problema de interpolación.

Teorema 1.18. ⁶(Schoenberg) *Una función ϕ es completamente monótona pero no constante en $[0, \infty)$ si, y solo si, $\Phi = \phi(\|\cdot\|_2^2)$ es definida positiva y radial en $\mathbb{R}^d$ para todo d .*

Observe que la función Φ ahora se define mediante el cuadrado de la norma.

Ejemplo 1.19. *Ejemplos de funciones completamente monótonas pero no constantes son:*

1. *La función Gaussiana*

$$\phi(r) = e^{-\epsilon r}, \quad \epsilon > 0,$$

es completamente monótona en $[0, \infty)$, y es no constante ahora por el Teorema 1.18 la función gaussiana $\Phi(x) = \phi(\|x\|_2^2) = e^{-\epsilon\|x\|_2^2}$, $\epsilon \geq 0$, es definida positiva y radial en $\mathbb{R}^d$ para todo d .

2. *La función multicuadrática inversa*

$$\phi(r) = (r + c^2)^{-\beta} \quad c, \beta > 0,$$

es completamente monótona en $[0, \infty)$, y es no constante ahora por el Teorema 1.18 la función multicuadrática inversa $\Phi(x) = \phi(\|x\|_2^2) = \frac{1}{(c^2 + \|x\|_2^2)^\beta}$, $c > 0$, $\beta \geq 0$, es definida positiva y radial en $\mathbb{R}^d$ para todo d .

Con base en lo anterior, los interpolantes (2) correspondientes toman la forma

$$S(x) = \sum_{j=1}^N \lambda_j e^{-\epsilon \Phi(\|x - x_j\|_2^2)}, \quad y \quad S(x) = \sum_{j=1}^N \lambda_j \Phi(\|x - x_j\|_2^2 + c^2)^{-\beta},$$

respectivamente.

1.3. FUNCIONES CONDICIONALMENTE DEFINIDAS POSITIVAS

La siguiente sección generaliza la noción de funciones definidas positivas de manera que se mejoren algunas propiedades relevantes para las funciones radiales base, como aquellas en las que el proceso de interpolación requiere de la inclusión de un polinomio. En consecuencia, se da una noción de la reproducción polinómica y su relación con las funciones condicionalmente semi definidas positivas y las funciones condicionalmente definidas positivas de orden m .

Definición 1.20. Una función continua $\Phi : \mathbb{R}^d \rightarrow \mathbb{R}$ es llamada condicionalmente semi definida positiva de orden m en $\mathbb{R}^d$, si

$$\sum_{i=1}^N \sum_{j=1}^N \alpha_i \alpha_j \Phi(x_i - x_j) \geq 0, \quad (17)$$

para cualesquiera N puntos $x_1, \dots, x_N \in \mathbb{R}^d$ distintos dos a dos, y $\alpha = [\alpha_1, \dots, \alpha_N]^T \in \mathbb{R}^N$ satisfaciendo

$$\sum_{j=1}^N \alpha_j p(x_j) = 0, \quad \forall p \in \mathbb{P}_{m-1}(\mathbb{R}^d). \quad (18)$$

La función Φ es condicionalmente definida positiva de orden m en $\mathbb{R}^d$ si la forma cuadrática (17) es cero solo para $\alpha = 0$.

Proposición 1.21. ⁶ Una función que es condicionalmente definida positiva de orden m en $\mathbb{R}^d$ es también condicionalmente definida positiva de cualquier orden superior. En particular, una función definida positiva es siempre condicionalmente definida positiva de cualquier orden.

Ejemplo 1.22. Ejemplos de funciones que son condicionalmente definidas positivas pero no definidas positivas son dadas por la función multicuadrática $\Phi(x) = \sqrt{c^2 + \|x\|_2^2}$, $c > 0$, y la función spline de placa delgada $\Phi(x) = \|x\|^2 \log(\|x\|)$.

A partir de la proposición anterior, toda función definida positiva es también condicionalmente definida positiva de cualquier orden. Es natural entonces, buscar el orden m más pequeño posible. Por tanto, cuando se habla de una función condicionalmente definida positiva de orden m , damos en general el mínimo posible de tales m .

Como se mencionó, las funciones condicionalmente definidas positivas son utilizadas para resolver el problema de interpolación (1) a partir de funciones radiales que requieren la inclusión de un polinomio. Más exactamente, permite formar interpolantes de la forma

$$s(x) = \sum_{j=1}^N \lambda_j \Phi(x - x_j) + \sum_{l=1}^Q \lambda_l p_l(x), \quad (19)$$

donde $Q = \dim(\mathbb{P}_{m-1}(\mathbb{R}^d))$ y $p_1, \dots, p_Q$ forman una base en el espacio lineal $\mathbb{P}_{m-1}(\mathbb{R}^d)$. Los coeficientes en (19) son definidos como

$$s(x_i) = f(x_i), \quad i = 1, \dots, N, \quad (20)$$

$$\sum_{j=1}^N \lambda_j p_l(x_j) = 0, \quad l = 1, \dots, Q. \quad (21)$$

Resolver el problema anterior es equivalente a resolver el sistema matricial

$$\begin{bmatrix} A & P \\ P^T & 0 \end{bmatrix} \begin{bmatrix} \lambda_j \\ \lambda_l \end{bmatrix} = \begin{bmatrix} f \\ 0 \end{bmatrix}, \quad (22)$$

donde A y P son matrices de orden $N \times N$ y $N \times Q$ respectivamente, con componentes dadas por

$$A = A_{i,j} = \Phi(x_i - x_j), \quad P = p_l(x_j). \quad (23)$$

Para determinar la no singularidad de (22) primero definimos la unisolvencia de un conjunto de puntos en $\mathbb{R}^d$, la cual se debe satisfacer dependiendo del grado del polinomio.

Definición 1.23. *El conjunto de puntos $X = \{x_1, x_2, \dots, x_N\} \subset \mathbb{R}^d$ con $N \geq Q = \dim(\mathbb{P}_m(\mathbb{R}^d))$ es llamado $\mathbb{P}_m(\mathbb{R}^d)$ -unisolvente si el polinomio cero es el único polinomio de $\mathbb{P}_m(\mathbb{R}^d)$ que se anula en todos ellos.*

Ejemplo 1.24. *Considere los polinomios lineales en $\mathbb{R}^2$. Sabemos que $\dim(\mathbb{P}_1(\mathbb{R}^2)) = 3$. Dado que cada polinomio lineal bivariado describe un plano en un espacio tridimensional, este plano está determinado únicamente por tres puntos si, y solo si, estos tres puntos no son colineales. Por lo tanto, el conjunto con tres puntos en $\mathbb{R}^2$ es $\mathbb{P}_1(\mathbb{R}^2)$ -unisolvente si, y solo si, los puntos no son colineales.*

Teorema 1.25. ⁶ *Supongamos que $\Phi: \rightarrow \mathbb{R}$ es una función condicionalmente definida positiva de orden m y X es un conjunto de centros (asociados a Φ y definidos en (2)) siendo $\mathbb{P}_{m-1}(\mathbb{R}^d)$ -unisolvente. Entonces el sistema (22) tiene una única solución.*

Ejemplo 1.26. *Un ejemplo de estas funciones son las spline de placa delgada, $\phi(r) = r^2 \log(r)$ con $r = \|\bar{x}\|$, $\bar{x} = (x, y) \in \mathbb{R}^2$. El polinomio requerido $p \in \mathbb{P}_1^2$ es bivariado lineal. Se tiene que la base del espacio polinomial $\mathbb{P}_1^2$ es $\{1, x, y\}$.*

En consecuencia, el interpolante tiene la forma

$$s(\bar{x}) = \sum_{j=1}^N \lambda_j \|\bar{x} - \bar{x}_j\|^2 \log(\|\bar{x} - \bar{x}_j\|) + \mu_1 + \mu_2 x + \mu_3 y, \quad \bar{x}, \bar{x}_j \in \mathbb{R}^2.$$

junto con las siguientes restricciones

$$\sum_{j=1}^N \lambda_j = 0, \quad \sum_{j=1}^N \lambda_j x_j = 0, \quad \sum_{j=1}^N \lambda_j y_j = 0.$$

En 1986 Micchelli ⁹ conjetura un criterio análogo al Teorema 1.18 de Schoenberg para el caso de funciones condicionalmente semi definidas positivas. Los siguientes teoremas reúnen los principales resultados de los trabajos citados.

Teorema 1.27. ⁷ *Supongamos que $\phi \in C[0, \infty) \cap C^\infty(0, \infty)$. La función $\Phi = \phi(\|\cdot\|_2^2)$ es condicionalmente semi definida positiva de orden $m \in \mathbb{N}_0$ y radial en $\mathbb{R}^d$ para toda d , si y solo si $(-1)^m \phi^{(m)}$ es completamente monótona en $(0, \infty)$.*

Además si ϕ no es un polinomio de grado máximo m , entonces Φ es condicionalmente definida positiva de orden m .

Ejemplo 1.28. *Las funciones multicuadráticas $\phi(r) = (-1)^{\lceil \beta \rceil} (c^2 + r^2)^\beta$, $c, \beta > 0$, $\beta \notin \mathbb{N}$, son condicionalmente definidas positivas de orden $m = \lceil \beta \rceil$ en todo.*

En efecto, definiendo $f_\beta(r) = (-1)^{\lceil \beta \rceil} (c^2 + r)^{\beta}$, se obtiene que

$$f_\beta^{(l)}(r) = (-1)^{\lceil \beta \rceil} \beta(\beta - 1) \cdots (\beta - l + 1) (c^2 + r)^{\beta - l},$$

⁹ C.A Micchelli. "Interpolation of scattered data-distance matrices and conditionally positive definite functions," en: *Constructive Approximation*, 2(1) (1986), págs. 11-22.

por lo que

$$(-1)^{\lceil \beta \rceil} f_{\beta}^{(\lceil \beta \rceil)}(r) = \beta(\beta - 1) \cdots (\beta - \lceil \beta \rceil + 1)(c^2 + r)^{\beta - \lceil \beta \rceil}$$

es completamente monótona. Más aún, $m = \lceil \beta \rceil$ es el número más pequeño posible tal que $(-1)^m f_{\beta}^{(m)}$ es completamente monótona. Como $\beta \notin \mathbb{N}$, f no es un polinomio, por lo tanto, $\Phi = \phi(\|\cdot\|)$ es condicionalmente definida positiva de orden $\lceil \beta \rceil$ y radial en $\mathbb{R}^d$ para todo d .

Dado el Teorema 1.27 se puede establecer la Tabla 1.1 con la generalización de las funciones condicionalmente definidas positivas más comunes.

Tabla 1. Funciones de base radial comúnmente usadas y su respectivo grado del polinomio.

Nombre	$\phi(r)$	Parámetros	Grado del polinomio
Multicuadrática	$(-1)^{\lceil \beta \rceil} (r^2 + c^2)^{\beta}$	$\beta > 0, \beta \notin \mathbb{N}$	$\lceil \beta \rceil$
Placa Delgada	$(-1)^{1+\beta/2} r^{\beta} \log r$	$\beta > 0, \beta \in 2\mathbb{N}$	$1 + \beta/2$
Potencias	$(-1)^{\lceil \beta/2 \rceil} r^{\beta}$	$\beta > 0, \beta \notin 2\mathbb{N}$	$\lceil \beta/2 \rceil$
Gaussiana	$e^{-\epsilon r^2}$	$\epsilon > 0$	0
Multicuadrática inverso	$(r^2 + c^2)^{-\beta/2}$	$\beta > 0, \beta \notin 2\mathbb{N}$	0

Se tiene que la función radial multicuadrática $\phi(r) = (c^2 + r^2)^{1/2}$ es una función condicionalmente definida positiva de orden uno, por tanto, se requiere de la adición de una constante para garantizar la invertibilidad del sistema lineal de ecuaciones. Sin embargo, Micchelli ⁹ demuestra que para la interpolación con funciones condicionalmente definidas positivas de orden uno, no es necesario incluir el polinomio constante.

Teorema 1.29. *Suponga que ϕ es condicionalmente definida positiva de orden 1 y que $\phi(0) \leq 0$. Para cualesquiera puntos distintos $x_1, \dots, x_N \in \mathbb{R}^d$ la matriz A con entradas $A_{i,j} = \phi(x_i - x_j)$ tiene $N - 1$ valores propios y 1 negativo, y por lo tanto es no singular.*

El Teorema 1.29 generaliza el teorema de Schoenberg para el caso $m = 1$. De modo que, cuando se trabaja con el núcleo multicuadrático es prescindible adicionar dicha

constante; numéricamente se ha observado que esto no afecta la exactitud en la interpolación.

1.4. ESPACIOS NATIVOS

Hasta ahora hemos encontrado funciones definidas positivas y condicionalmente definidas positivas en el contexto de un problema de interpolación de datos dispersos en $\mathbb{R}^d$. En esta sección veremos, desde otro ángulo, las funciones previamente mencionadas con el fin de presentar algunas ideas sobre estimación de errores en el proceso de interpolación.

Dentro de la teoría de la interpolación con funciones de base radial, los límites de error son estimados asociando a cada función básica radial, un cierto espacio de funciones llamado espacio nativo. Estos espacios están conectados con los llamados núcleos reproductores sobre espacios de Hilbert. Aunque un desarrollo completo y autocontenido de la teoría de los núcleos reproductores está fuera del alcance de este trabajo, presentamos algunos ingredientes necesarios para comprender mejor las estimaciones de error.

A continuación se describe la construcción de espacios nativos en base a la definición de núcleos reproductores en un espacio de Hilbert, esto es, dado un núcleo radial, construir su espacio de Hilbert asociado. Lo anterior es requerido para definir la convergencia del interpolante. Para ello, asumiremos la existencia de un núcleo Φ y se construirá el núcleo reproductor en espacios de Hilbert. Primero se mostrará la conexión entre espacios de Hilbert con núcleo reproductor y núcleos definidos positivos. Esta sección está fundamentada en ⁷.

Definición 1.30. Sea $\mathcal{H}$ un espacio de Hilbert de funciones $f : \Omega \rightarrow \mathbb{R}$ con producto interno $\langle \cdot, \cdot \rangle_{\mathcal{H}}$, siendo $\Omega \subset \mathbb{R}^d$ no vacío. Una función $\Phi : \Omega \times \Omega \rightarrow \mathbb{R}$ es llamada núcleo reproductor para $\mathcal{H}$ si

1. $\Phi(x, \cdot) \in \mathcal{H} \quad \forall x \in \Omega$,
2. $f(x) = \langle f, \Phi(\cdot, x) \rangle_{\mathcal{H}}, \quad \forall f \in \mathcal{H} \text{ y } \forall x \in \Omega$.

La propiedad 2 es llamada propiedad reproductora y se sigue fácilmente que si un espacio de Hilbert tiene un núcleo reproductor, este es único. En efecto, si Φ_1 y Φ_2 son núcleos reproductores, entonces por la propiedad 2 tenemos $\langle f, \Phi_1(\cdot, y) - \Phi_2(\cdot, y) \rangle_{\mathcal{H}} = 0$ para todo $f \in \mathcal{H}$ y $y \in \Omega$. Tomando $f = \Phi_1(\cdot, y) - \Phi_2(\cdot, y)$ en la última igualdad, se tiene que $\langle \Phi_1(\cdot, y) - \Phi_2(\cdot, y), \Phi_1(\cdot, y) - \Phi_2(\cdot, y) \rangle_{\mathcal{H}} = \|\Phi_1(\cdot, y) - \Phi_2(\cdot, y)\|_{\mathcal{H}} = 0$, y como $y \in \Omega$ es arbitrario, se concluye que $\Phi_1 = \Phi_2$.

Una caracterización de los espacios de Hilbert que tienen un núcleo reproductor es dada por la siguiente proposición.

Proposición 1.31. ⁷ *Supongamos que $\mathcal{H}$ es un espacio de Hilbert compuesto de funciones $f : \Omega \rightarrow \mathbb{R}$. Entonces son equivalentes:*

1. *El espacio $\mathcal{H}$ tiene núcleo reproductor.*
2. *Los funcionales de evaluación $\delta_x : f \mapsto f(x)$, $f \in \mathcal{H}$, son continuos para todo $x \in \Omega$, es decir, $\delta_x \in \mathcal{H}'$ para cada $x \in \Omega$.*

Demostración. Supongamos que el espacio $\mathcal{H}$ tiene núcleo reproductor Φ . Entonces $\delta_x = \langle \cdot, \Phi(\cdot, x) \rangle_{\mathcal{H}}$, para cada $x \in \Omega$. Dado que el producto interno es continuo. Por otro lado, si los funcionales de evaluación son continuos, por el Teorema de representación de Riesz, para cada $y \in \Omega$ existe $\Phi_y \in \mathcal{H}$ tal que $\delta_y(f) = \langle f, \Phi_y \rangle$ para todo $f \in \mathcal{H}$. Así, $\Phi : \Omega \times \Omega \rightarrow \mathbb{R}$ definido por $\Phi(x, y) := \Phi_y(x)$ es el núcleo reproductor de $\mathcal{H}$. $\square$

El siguiente teorema enuncia algunas propiedades importantes de un espacio de Hilbert con núcleo reproductor.

Teorema 1.32. ⁷ *Supongamos que $\mathcal{H}$ es un espacio de Hilbert compuesto de funciones $f : \Omega \rightarrow \mathbb{R}$, con núcleo reproductor Φ . Entonces*

1. *$\Phi(x, y) = \langle \Phi(\cdot, x), \Phi(\cdot, y) \rangle_{\mathcal{H}}$, para $x, y \in \Omega$.*
2. *$\Phi(x, y) = \Phi(y, x)$, para $x, y \in \Omega$.*
3. *Si f, f_n , $n \in \mathbb{N}$, son funciones tales que f_n converge a f en $\mathcal{H}$, entonces f_n converge a f puntualmente.*

El siguiente resultado muestra la conexión entre espacios de Hilbert con núcleo reproductor y núcleos definidos positivos.

Teorema 1.33. ⁷ *Supongamos que $\mathcal{H}$ es un espacio de Hilbert con núcleo reproductor $\Phi : \Omega \times \Omega \rightarrow \mathbb{R}$. Entonces Φ es semi definido positivo. Además, Φ es definido positivo si, y solo si, los funcionales de evaluación son linealmente independientes en $\mathcal{H}'$.*

Hasta ahora hemos visto que el núcleo reproductor de un espacio de funciones de Hilbert origina una función semi definida positiva; sin embargo, normalmente no comenzamos con un espacio de funciones, sino con un núcleo semi definido positivo, con lo cual nos enfrentamos al problema de encontrar el espacio de funciones asociado que tiene este núcleo como su núcleo reproductor. A continuación queremos discutir brevemente este problema construyendo el espacio de funciones de Hilbert correspondiente. Queremos mostrar que toda función básica radial definida positiva puede asociarse con un núcleo reproductor sobre un espacio de Hilbert, a saber, su espacio nativo. Para ello vamos a suponer que $\Phi : \Omega \times \Omega \rightarrow \mathbb{R}$ es un núcleo simétrico semi definido positivo.

De la primera propiedad de la Definición 1.30 se deduce que $\mathcal{H}$ contiene todas las funciones de la forma:

$$f(\cdot) = \sum_{j=1}^N c_j \Phi(\cdot, x_j), \quad (24)$$

desde que $x_j \in \Omega$. Además, para $f \in \mathcal{H}$, se tiene que,

$$\|f\|_{\mathcal{H}} = \langle f, f \rangle_{\mathcal{H}} = \left\langle \sum_{j=1}^N c_j \Phi(\cdot, x_j), \sum_{k=1}^N c_k \Phi(\cdot, x_k) \right\rangle = \sum_{j=1}^N \sum_{k=1}^N c_j c_k \Phi(x_j, x_k). \quad (25)$$

Se define el espacio $H_{\Phi}(\Omega)$ como el espacio lineal generado por todas las funciones de la forma (24)

$$H_{\Phi}(\Omega) = \text{gen}\{\Phi(\cdot, x) : x \in \Omega\}, \quad (26)$$

el cual tiene una forma bilineal asociada, definida de la siguiente manera: Para cua-

lesquiera par de funciones $s, u \in H_\Phi(\Omega)$,

$$s = \sum_{j=1}^N c_j \Phi(\cdot, x_j), \quad u = \sum_{k=1}^N d_k \Phi(\cdot, y_k),$$

la forma bilineal asociada se define como

$$\langle u, s \rangle_\Phi = \langle s, u \rangle_\Phi = \sum_{j=1}^N \sum_{k=1}^N c_j d_k \Phi(x_j, y_k), \quad (27)$$

Teorema 1.34. ⁷ Sea $\Phi : \Omega \times \Omega \rightarrow \mathbb{R}$, $\Omega \subset \mathbb{R}^d$ un núcleo simétrico y definido positivo. Entonces la forma bilineal $\langle \cdot, \cdot \rangle_\Phi$ define un producto interior en $H_\Phi(\Omega)$. Adicionalmente, $H_\Phi(\Omega)$ es un espacio pre-Hilbert que tiene como núcleo reproductor a Φ .

Demostración. Claramente $\langle \cdot, \cdot \rangle_\Phi$ es bilineal y simétrico, puesto que Φ es simétrico. Resta mostrar que $\langle f, f \rangle_\Phi > 0$ para $f \in H_\Phi(\Omega)$ no nulo. En efecto, para cualquier $f \in H_\Phi(\Omega)$ no nulo, este se puede escribir como

$$f = \sum_{j=1}^N c_j \Phi(\cdot, x_j), \quad x_j \in \Omega.$$

Entonces

$$\langle f, f \rangle_\Phi = \sum_{j=1}^N \sum_{k=1}^N c_j c_k \Phi(x_j, x_k) > 0,$$

ya que Φ es definido positivo. Finalmente, dado que

$$\langle f, \Phi(\cdot, x) \rangle_\Phi = \sum_{j=1}^N c_j \Phi(x_j, x) = f(x),$$

entonces se verifica la propiedad de reproducción. □

Debido a que el teorema anterior establece que $H_\Phi(\Omega)$ es un espacio pre-Hilbert, es decir, no es necesariamente completo, el primer candidato para un espacio de Hilbert con núcleo reproductor Φ es la completación $\mathcal{H}_\Phi$ de $H_\Phi(\Omega)$ respecto a la norma $\| \cdot \|_\Phi$. Sin embargo, los elementos de la completación quedan simplemente como

elementos abstractos de un espacio de Hilbert. Para representar a estos elementos como funciones, consideremos primero la aplicación lineal

$$\begin{aligned} R : \mathcal{H}_\Phi(\Omega) &\rightarrow C(\Omega) \\ f &\mapsto \langle f, \Phi(\cdot, x) \rangle_\Phi. \end{aligned} \quad (28)$$

Podemos ver que la aplicación R está bien definida; en efecto, por la desigualdad de Cauchy - Schwarz tenemos,

$$|Rf(x) - Rf(y)| = |\langle f, \Phi(\cdot, x) - \Phi(\cdot, y) \rangle_\Phi| \leq \|f\|_\Phi \|\Phi(\cdot, x) - \Phi(\cdot, y)\|_\Phi, \quad (29)$$

dado que

$$\|\Phi(\cdot, x) - \Phi(\cdot, y)\|_\Phi^2 = \Phi(x, x) + \Phi(y, y) - 2\Phi(x, y), \quad (30)$$

entonces $Rf(x) \rightarrow Rf(y)$ cuando $x \rightarrow y$ ya que Φ es continua, por tanto, $R(f) \in C(\Omega)$. Más aún se tiene que $Rf(x) = f(x)$ para todo $x \in \Omega$ y todo $f \in H_\Phi(\Omega)$.

Proposición 1.35. *La aplicación $R : \mathcal{H}_\Phi(\Omega) \rightarrow C(\Omega)$ definida en (28) es inyectiva.*

Demostración. Supongamos que $Rf(x) = 0$ para cualquier $f \in \mathcal{H}_\Phi(\Omega)$, es decir, $\langle f, \Phi(\cdot, x) \rangle_\Phi = 0$ para todo $x \in \Omega$. Pero $\mathcal{H}_\Phi$ es la completación de $H_\Phi(\Omega)$. Por lo tanto, el único elemento de $\mathcal{H}_\Phi(\Omega)$ perpendicular a $H_\Phi(\Omega)$ es $f = 0$. $\square$

Como la aplicación es inyectiva y $Rf(x) = f(x)$ para toda $x \in \Omega$ y $f \in H_\Phi(\Omega)$ podemos dar la definición de un espacio nativo.

Definición 1.36. *El espacio nativo correspondiente al núcleo simétrico y definido positivo $\Phi : \Omega \times \Omega \rightarrow \mathbb{R}$ se define como*

$$\mathcal{N}_\Phi(\Omega) := R(\mathcal{H}_\Phi(\Omega)), \quad (31)$$

con el producto interior

$$\langle f, g \rangle_{\mathcal{N}_\Phi(\Omega)} := \langle R^{-1}f, R^{-1}g \rangle_\Phi.$$

El espacio así definido es un espacio de Hilbert de funciones continuas con núcleo de reproducción Φ . Dado que $\Phi(\cdot, x)$ es un elemento de $H_\Phi(\Omega)$ para $x \in \Omega$, Φ no cambia bajo la aplicación lineal R , por lo tanto,

$$f(x) = \langle R^{-1}(f), \Phi(\cdot, x) \rangle_\Phi = \langle f, \Phi(\cdot, x) \rangle_{\mathcal{N}_\Phi(\Omega)}, \quad (32)$$

para todo $f \in \mathcal{N}_\Phi(\Omega)$ y $x \in \Omega$.

Teorema 1.37. ⁷ Si $\Phi : \Omega \times \Omega \rightarrow \mathbb{R}$ es un núcleo simétrico y definido positivo, entonces su espacio nativo asociado $\mathcal{N}_\Phi(\Omega)$ es un espacio de Hilbert con núcleo reproductor Φ . Este espacio es además único.

Por lo tanto, los núcleos definidos positivos y los núcleos reproductores de espacios de funciones de Hilbert son lo mismo.

Finalizamos esta sección determinando un espacio nativo, estudiado por Madych y Nelson en ⁶ para las funciones definidas positivas para el caso especial $\mathbb{R}^d$. Esta caracterización del espacio nativo está expresada en términos de la transformada de Fourier.

Teorema 1.38. Suponga que $\Phi \in C(\mathbb{R}^d) \cap L^1(\mathbb{R}^d)$ es una función definida positiva que toma valores en los reales. Definimos

$$\mathcal{G} := \left\{ f \in L^2(\mathbb{R}^d) \cap C(\mathbb{R}^d) : \frac{\hat{f}}{\sqrt{\hat{\Phi}}} \in L^2(\mathbb{R}^d) \right\}$$

para la cual su forma bilineal esta dada por

$$\langle f, g \rangle_{\mathcal{G}} := (2\pi)^{-d/2} \left\langle \frac{\hat{f}}{\sqrt{\hat{\Phi}}}, \frac{\hat{g}}{\sqrt{\hat{\Phi}}} \right\rangle_{L^2(\mathbb{R}^d)} = (2\pi)^{-d/2} \int_{\mathbb{R}^d} \frac{\hat{f}(w) \overline{\hat{g}(w)}}{\hat{\Phi}(w)} dw.$$

Entonces $\mathcal{G}$ es un espacio de Hilbert real con producto interno $\langle \cdot, \cdot \rangle_{\mathcal{G}(\mathbb{R}^d)}$ y núcleo reproductor $\Phi(\cdot - \cdot)$. Por lo tanto, $\mathcal{G}$ es el espacio nativo de Φ en $\mathbb{R}^d$ con la norma canónica inducida por el producto interior, es decir, $\mathcal{G} = \mathcal{N}_\Phi(\mathbb{R}^d)$.

Ejemplo 1.39. Recordemos que el espacio de Sobolev de orden s para $s > d/2$ se

puede definir como $H^s(\mathbb{R}^d) = \{u \in L^2(\mathbb{R}^d) \cap C(\mathbb{R}^d) : \hat{u}(\cdot)(1 + \|\cdot\|^2)^{s/2} \in L^2(\mathbb{R}^d)\}$. Supongamos ahora que $\Phi \in L^1(\mathbb{R}^d) \cap C(\mathbb{R}^d)$ satisface la siguiente condición de decaimiento algebraico

$$c_1(1 + \|\omega\|^2)^s \leq \hat{\Phi}(\omega) \leq c_2(1 + \|\omega\|^2)^s, \quad \omega \in \mathbb{R}^d,$$

con $s > d/2$ y $c_1 \leq c_2$ positivos. Entonces, por el Teorema 1.38, el espacio nativo $\mathcal{N}_\Phi(\mathbb{R}^d)$ asociado a Φ corresponde con el espacio de Sobolev $H^s(\mathbb{R}^d)$ y las normas son equivalentes. En este sentido, podemos pensar en los espacios nativos como una cierta generalización de los espacios de Sobolev.

1.5. ESTIMACIONES DE ERROR

El objetivo de esta sección es proporcionar algunas estimaciones de error para la interpolación de datos dispersos con funciones definidas positivas. La siguiente sección ha sido desarrollada basándonos en los trabajos de Duchon ¹⁰, Fassahuer⁶, Madych y Nelson ¹¹ y Wu-Schaback ¹². Se quiere que este tipo de estimaciones dependan de algún tipo de medida de la distribución de datos, con lo cual se introduce una función que generalmente se usa en la teoría de la aproximación denominada *distancia de llenado* asociada a un conjunto de datos (nodos) $X = \{x_1, \dots, x_N\}$ distribuidos de manera no uniforme. La función de llenado se define como

$$h = h_{X,\Omega} = \sup_{x \in \Omega} \min_{x_i \in X} \|x - x_i\|_2,$$

la cual, como su nombre lo indica, mide la manera como los datos llenan el dominio, y corresponde a la distancia máxima de cualquier punto del dominio a su punto de

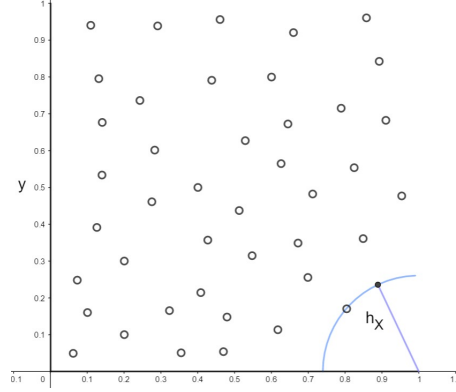
¹⁰ J. Duchon. "Sur l'erreur d'interpolation des fonctions de plusieurs variables par les D^m -splines," en: *Rev. Française Automat. Informat. Rech. Opér. Anal.* 12 (1978), págs. 325-334.

¹¹ Nelson S. A. Madych W. R. "Polyharmonic cardinal splines," en: *Journal of Approximation Theory*, 60 (1990), págs. 141-156.

¹² Z. Wu y R. Schaback. "Local error estimates for radial basis function interpolation of scattered data," en: *IMA Journal of Numerical Analysis*, 13(1) (1993), págs. 13-21.

referencia más cercano, o equivalentemente, el radio de la bola vacía más grande que puede situarse entre los datos (ver Figura 3).

Figura 3. Distancia de llenado



Estamos interesados en saber si el error $\|f - \mathcal{P}_f^{(h)}\|_\infty$ converge a cero cuando h tiende a cero, y qué tan rápido. Aquí, $\{\mathcal{P}^{(h)}\}_h$ denota una sucesión de operadores de interpolación, la cual depende de la distancia de llenado. Diremos que el proceso de aproximación es de orden k si $\|f - \mathcal{P}_f^{(h)}\|_q = \mathcal{O}(h^k)$, $h \rightarrow 0$, donde $\|\cdot\|_q$ denota la norma en L^q , $1 \leq q \leq \infty$. Dado que queremos emplear la teoría de núcleos reproductores en espacios de Hilbert, nos concentraremos en las estimaciones de error para funciones $f \in \mathcal{N}_\Phi$.

La idea ahora es expresar al interpolante en su forma de Lagrange, utilizando las llamadas funciones cardinales. Esta idea se debe a Wu y Schaback ¹². Como se vió en la Sección 1.2, para cualquier función definida positiva Φ , el sistema $A\lambda = y$, con $A_{ij} = \Phi(x_i - x_j)$, $i, j = 1, \dots, N$, $\lambda = [\lambda_1, \dots, \lambda_N]^T$, y $y = [f(x_1), \dots, f(x_N)]^T$, tiene una única solución. A continuación consideraremos la situación más general donde Φ es un núcleo definido positivo, es decir, las entradas de A están dados por $A_{ij} = \Phi(x_i, x_j)$. El resultado de unicidad también es válido en este caso. Para obtener las funciones de la base cardinal de Lagrange u_j^* , $1 \leq j, i \leq N$, con la propiedad de que $u_j^*(x_i) = \delta_{ij}$, es decir,

$$u_j^*(x_i) = \begin{cases} 1 & \text{si } j = i, \\ 0 & \text{si } j \neq i, \end{cases}$$

consideramos el sistema lineal

$$Au^*(x) = b(x),$$

donde $u^* = [u_1^*, \dots, u_N^*]^T$ y $b = [\Phi(\cdot, x_1), \dots, \Phi(\cdot, x_N)]^T$. Entonces se tiene el siguiente resultado.

Teorema 1.40. ⁷ *Suponga que Φ es un núcleo definido positivo sobre $\mathbb{R}^d$. Entonces, para cualquier conjunto de puntos distintos $x_1, \dots, x_N$, existen funciones $u_j^* \in \text{gen}\{\Phi(\cdot, x_j), j = 1, \dots, N\}$ tales que $u_j^*(x_i) = \delta_{ij}$.*

Teniendo en cuenta lo anterior, podemos escribir el interpolante $\mathcal{P}_f$ de f en $x_1, \dots, x_N$ como

$$\mathcal{P}_f(x) = \sum_{j=1}^N f(x_j)u_j^*(x), \quad x \in \mathbb{R}^d. \quad (33)$$

Otro elemento que requerimos para determinar las cotas de error entre el interpolante y los datos es la llamada función potencia. Para ello consideremos que $\Phi \in C(\Omega \times \Omega)$ con $\Omega \subseteq \mathbb{R}^d$, y para cualquier conjunto de puntos distintos $X = \{x_1, \dots, x_N\}$ contenidos en Ω , y para cualquier vector $u = \{u_1, \dots, u_N\} \in \mathbb{R}^N$, definimos la forma cuadrática

$$Q(u) = \Phi(x, x) - 2 \sum_{j=1}^N u_j \Phi(x, x_j) + \sum_{i=1}^N \sum_{j=1}^N u_i u_j \Phi(x_i, x_j), \quad (34)$$

la cual es equivalente a la norma inducida por el producto interno en el espacio nativo $\mathcal{N}_{\Phi(\Omega)}$

$$Q(u) = \|\Phi(\cdot, x) - \sum_{j=1}^N u_j \Phi(\cdot, x_j)\|_{\mathcal{N}_{\Phi(\Omega)}}^2. \quad (35)$$

Ahora podemos definir la función potencia como:

Definición 1.41. *Sea $\Omega \subseteq \mathbb{R}^d$ y suponga que $\Phi \in C(\Omega \times \Omega)$ una función definida positiva en $\mathbb{R}^d$. Para cualesquiera puntos distintos $X = \{x_1, \dots, x_N\} \subseteq \Omega$, la función potencia está definida por*

$$[P_{\Phi, X}(x)]^2 = Q(u^*(x)),$$

donde u^* es el vector de funciones cardinales.

La cota de error genérica entre el interpolante y los datos de la función está dada en términos del siguiente resultado

Teorema 1.42. Sea $\Omega \subseteq \mathbb{R}^d$, $\Phi \in C(\Omega \times \Omega)$ una función definida positiva en $\mathbb{R}^d$, y suponga que el conjunto de puntos $X = \{x_1, \dots, x_N\} \in \Omega$ son todos puntos distintos. Denote $s(x)$ al interpolante de $f \in \mathcal{N}_{\Phi(\Omega)}$ en X , entonces para cualquier $x \in \Omega$

$$|f(x) - s(x)| \leq P_{\Phi, X}(x) \|f\|_{\mathcal{N}_{\Phi(\Omega)}}. \quad (36)$$

Demostración. Supongamos que f se encuentra en el espacio nativo de Φ , por la propiedad de reproducción de Φ se tiene que:

$$f(x) = \langle f, \Phi(\cdot, x) \rangle_{\mathcal{N}_{\Phi(\Omega)}}.$$

Expresamos el interpolante en su forma cardinal y aplicamos la propiedad de reproducción de Φ . Esto explica que

$$\begin{aligned} s(x) &= \sum_{j=1}^N f(x_j) u_j^*(x) \\ &= \sum_{j=1}^N u_j^*(x) \langle f, \Phi(\cdot, x_j) \rangle_{\mathcal{N}_{\Phi(\Omega)}} \\ &= \left\langle f, \sum_{j=1}^N u_j^*(x) \Phi(\cdot, x_j) \right\rangle_{\mathcal{N}_{\Phi(\Omega)}}. \end{aligned}$$

Ahora todo lo que resta por hacer es combinar las dos fórmulas recién derivadas y aplicar la desigualdad de Cauchy-Schwarz. Así,

$$\begin{aligned} |f(x) - s(x)| &= \left| \langle f, \Phi(\cdot, x) - \sum_{j=1}^N u_j^*(x) \Phi(\cdot, x_j) \rangle_{\mathcal{N}_{\Phi(\Omega)}} \right| \\ &\leq \|f\|_{\mathcal{N}_{\Phi(\Omega)}} \left\| \Phi(\cdot, x) - \sum_{j=1}^N u_j^*(x) \Phi(\cdot, x_j) \right\|_{\mathcal{N}_{\Phi(\Omega)}} \\ &\leq \|f\|_{\mathcal{N}_{\Phi(\Omega)}} P_{\Phi, X}(x). \end{aligned}$$

□

Un beneficio de este tipo de estimación radica en que nos permite dividir el error entre la función desconocida f del espacio nativo y su interpolante, en dos términos, siendo

un término (la función de potencia) independiente de f y el otro independiente de la distribución de los centros X .

La estimación de error encontrada es análoga a la estimación del error en la interpolación polinomial ⁶, esto es,

$$|f(x) - If(x)| \leq P(x) \max_{x \in I} |f^n(x)|,$$

en donde la función $P_{\Phi, X} \simeq P(x)$ es de la forma

$$P(x) = \frac{1}{n!} \prod_{j=1}^n |x - x_j|.$$

Además, un cambio del parámetro de forma c de cualquier función básica Φ , tendrá como consecuencia alteraciones en la estimación de error dada en (36).

El siguiente paso es refinar la cota de error (36) para determinar cómo influye la distribución de los nodos X en términos de la distancia de llenado.

Teorema 1.43. *Sea $\Omega \subseteq \mathbb{R}^d$, $\Phi \in C(\Omega \times \Omega)$ es definida positiva en $\mathbb{R}^d$. Sea $X = \{x_1, x_2, \dots, x_N\} \subseteq \Omega$ un conjunto de puntos distintos, defina la forma cuadrática $Q(u)$ como en (35). El mínimo de $Q(u)$ es dado por el vector $u = u^*(x)$ del Teorema 1.40, es decir,*

$$Q(u^*(x)) \leq Q(u) \text{ para todo } u \in \mathbb{R}^N.$$

Para lograr este refinamiento es necesario proporcionar el siguiente teorema que nos garantiza la existencia de un vector $\tilde{u}$ que facilita la precisión polinomial deseada, es decir, el siguiente teorema asegura la existencia de un esquema de aproximación con precisión polinomial local. Este teorema requiere la noción de un dominio que satisface una condición del cono interior.

Definición 1.44. *Una región $\Omega \subseteq \mathbb{R}^d$ satisface la condición de cono interior si existe un ángulo $\theta \in (0, \pi/2)$ y un radio $r > 0$ tal que para todo $x \in \Omega$ existe un único vector $\xi(x)$*

tal que el cono

$$C = \{x + \lambda y : y \in \mathbb{R}^d, \|y\|_2 = 1, y^T \xi(x) \leq \cos(\theta), \lambda \in [0, r]\}$$

esta contenido en Ω .

Teorema 1.45. *Suponga que $\Omega \subseteq \mathbb{R}^d$ es acotado y satisface la condición de cono interior, y sea l un entero no negativo. Entonces existen constantes h_0, r_1 y r_2 tales que para todo $X = \{x_1, \dots, x_N\} \subseteq \Omega$ con $h_{X,\Omega} \leq h_0$ y todo $x \in \Omega$ existen números $\tilde{u}_1(x), \dots, \tilde{u}_N(x)$ con*

1. $\sum_{j=1}^N \tilde{u}_j(x) p(x_j) = p(x)$ para todo polinomio $p \in \mathbb{P}_l^d(\Omega)$,
2. $\sum_{j=1}^N |\tilde{u}_j(x)| \leq r_1$,
3. $\tilde{u}_j(x) = 0$ si $\|x - x_j\|_2 \geq r_2 h_{X,\Omega}$.

El siguiente resultado determina el error genérico refinado basado en la distancia de llenado.

Teorema 1.46. ⁶ *Suponga que $\Omega \subseteq \mathbb{R}^d$ es abierto y acotado y satisface la condición de cono interior. Suponga que $\Phi \in C^{2k}(\Omega \times \Omega)$ es simétrica y definida positiva en $\mathbb{R}^d$, y denote por s al interpolante de $f \in \mathcal{N}_{\Phi(\Omega)}$. Entonces existen constantes positivas h_0 y T (independientes de x, f y Φ) tales que*

$$|f(x) - s(x)| \leq T h_{X,\Phi}^k \sqrt{T_{\Phi}(x)} \|f\|_{\mathcal{N}_{\Phi(\Omega)}}, \quad x \in \Omega,$$

siempre que $h_{X,\Phi} \leq h_0$, donde

$$T_{\Phi}(x) = \max_{|\beta|=2k} \max_{w,z \in \Omega \cap B(x, r_2 h_{X,\Omega})} |D_2^{\beta} \Phi(w, z)|$$

con $B(x, r_2 h_{X,\Omega})$ siendo la bola de radio $r_2 h_{X,\Omega}$ y centro x .

Demostración. Por Teorema 1.42 se tiene que

$$|f(x) - s(x)| \leq P_{\Phi,X}(x) \|f\|_{\mathcal{N}_{\Phi(\Omega)}}.$$

Por lo tanto, ahora derivamos la cota

$$P_{\Phi, X}(x) \leq Th_{X, \Phi}^k \sqrt{T_{\Phi}(x)}$$

para la función de potencia en términos de la distancia de llenado. Sabemos que la función de potencia está definida por

$$[P_{\Phi, X}(x)]^2 = Q(u^*(x)).$$

Además, por el Teorema 1.43, la forma cuadrática $Q(u)$ está minimizada por $u = u^*(x)$. Por lo tanto, cualquier otro vector de coeficientes u producirá un límite superior en la función de potencia. Tomamos $u = \tilde{u}(x)$ del Teorema 1.45 para asegurarnos de tener precisión polinomial de grado $i > 2k - 1$. Para esta elección específica de coeficientes tenemos

$$[P_{\Phi, X}(x)]^2 \leq Q(u(x)) = \Phi(x, x) - 2 \sum_j u_j \Phi(x, x_j) + \sum_i \sum_j u_i u_j \Phi(x_i, x_j),$$

donde las sumas están sobre esos índices j con $u_j \neq 0$. Ahora aplicamos la expansión de Taylor centrada en x a $\Phi(x, \cdot)$ y centrada en x_i a $\Phi(x_i, \cdot)$, y evaluamos ambas funciones en x_j . Esto produce

$$\begin{aligned} Q(u) = \Phi(x, x) - 2 \sum_j u_j \left[\sum_{|\beta| < 2k} \frac{D_2^\beta \Phi(x, x)}{\beta!} (x_j - x)^\beta + R(x, x_j) \right] \\ + \sum_i \sum_j u_i u_j \left[\sum_{|\beta| < 2k} \frac{D_2^\beta \Phi(x_i, x_i)}{\beta!} (x_j - x_i)^\beta + R(x_i, x_j) \right], \end{aligned}$$

donde $D_2^\beta \Phi(x, x)$ denota el operador diferencial aplicado a $\Phi(x, x)$, $R(x, x_j)$ es el resto de la expansión de Taylor, definido como $R(x, x_j) = \sum_{|\beta| = 2k} \frac{D_2^\beta \Phi(x, \xi_{x, x_j})}{\beta!} (x_j - x)^\beta$, tal que ξ_{x, x_j} se encuentra en algún lugar del segmento de línea que conecta x y x_j . A continuación, identificamos $p(z) = (z - x)^\beta$ de modo que $p(x) = 0$ a menos que $\beta = 0$. Por lo tanto, la propiedad de precisión polinomial del vector de coeficientes u simplifica

esta expresión a

$$Q(u) = \Phi(x, x) - 2\Phi(x, x) - 2 \sum_j u_j R(x, x_j) + \sum_i u_i \sum_{|\beta| < 2k} \frac{D_2^\beta \Phi(x_i, x_i)}{\beta!} (x - x_i)^\beta + \sum_i \sum_j u_i u_j R(x_i, x_j). \quad (37)$$

Ahora podemos aplicar la expansión de Taylor nuevamente y hacer la observación de que

$$\sum_{|\beta| < 2k} \frac{D_2^\beta \Phi(x_i, x_i)}{\beta!} (x - x_i)^\beta = \Phi(x_i, x) - R(x_i, x). \quad (38)$$

De las ecuaciones (37) y (38) se tiene que

$$Q(u) = -\Phi(x, x) - \sum_j u_j \left[2R(x, x_j) - \sum_i u_i R(x_i, x_j) \right] + \sum_i u_i [\Phi(x_i, x) - R(x_i, x)]. \quad (39)$$

De la simetría de Φ ,

$$\Phi(x_i, x) = \Phi(x, x_i) = \sum_{|\beta| < 2k} \frac{D_2^\beta \Phi(x, x)}{\beta!} (x_i - x)^\beta + R(x, x_i). \quad (40)$$

Usando los resultados obtenidos en (39), (40) y la propiedad de precisión polinomial del vector de coeficientes u , nos quedamos con

$$Q(u) = - \sum_j u_j \left[R(x, x_j) + R(x_j, x) - \sum_i u_i R(x_i, x_j) \right].$$

Ahora el Teorema 1.45 nos permite acotar $\sum_j |u_j| \leq r_1$. Además, $\|x - x_j\| \leq r_2 h_{X, \Omega}$, y $\|x_i - x_j\|_2 \leq 2r_2 h_{X, \Omega}$. Por lo tanto, los tres términos restantes pueden acotarse por una expresión de la forma $Th_{X, \Omega}^{2k} T_\Phi(x)$. Aquí se usó la propiedad de cono interior de Ω permitiéndonos definir el término $T_\Phi(x)$. Combinando estos límites y tomando la raíz cuadrada se obtiene el límite establecido para la función de potencia. $\square$

El teorema anterior dice que la interpolación con un núcleo Φ suave de clase C^{2k} tiene

un orden de aproximación k . Por lo tanto, para funciones definidas positivas infinitamente suaves como las funciones gaussianas, y las multicuadráticas inversas generalizadas, vemos que el orden de aproximación k es arbitrariamente alto.

El enunciado del Teorema 1.46 puede generalizarse para funciones condicionalmente definidas positivas y también para cubrir el error de aproximar las derivadas de f por derivadas de s .

Teorema 1.47. *Sea $\Omega \subseteq \mathbb{R}^d$ es abierto y acotado y satisface la condición de cono interior. Suponga que $\Phi \in C^{2k}(\Omega \times \Omega)$ es simétrica y condicionalmente definida positiva de orden m en $\mathbb{R}^d$. Denote por s al interpolante de $f \in \mathcal{N}_{\Phi(\Omega)}$ para un conjunto $X = \{x_1, \dots, x_N\} \subset \Omega$ $\mathbb{P}_{m-1}(\mathbb{R}^d)$ -unisolvante. Para un valor fijo $\alpha \in \mathbb{N}_0^d$ con $|\alpha| \leq k$, existen h_0 y T constantes positivas (independiente de x, f , y Φ) tales que*

$$|D^\alpha f(x) - D^\alpha s(x)| \leq Th_{X,\Phi}^{k-|\alpha|} \sqrt{T_\Phi(x)} |f|_{\mathcal{N}_{\Phi(\Omega)}}, \quad x \in \Omega,$$

siempre que $h_{X,\Omega} \leq h_0$, donde

$$T_\Phi(x) = \max_{\substack{\beta, \gamma \in \mathbb{N}_0^d \\ |\beta| + |\gamma| = 2k}} \max_{w, z \in \Omega \cap B(x, r_2 h_{X,\Omega})} |D_1^\beta D_2^\gamma \Phi(w, z)|.$$

Es claro que el teorema anteriormente visto es genérico, debido a que no toma en cuenta explícitamente el núcleo Φ usado en la interpolación. Dependiendo del núcleo radial Φ , y bajo los supuestos anteriores es posible obtener una estimación del orden de convergencia en términos de la distancia de llenado y el núcleo radial seleccionado.

A continuación se establece una tasa de convergencia sobre el núcleo multicuadrático. El siguiente resultado muestra que, bajo ciertas condiciones, el orden de convergencia llega a ser exponencial.

Teorema 1.48. ⁷ *Sean $\Omega \subseteq \mathbb{R}^d$ un cubo, $\Phi = \phi(\|\cdot\|_2)$ una función radial condicionalmente definida positiva y $\psi = \phi(\sqrt{\cdot})$. Si existe l_0 entero no negativo tal que $|\psi^{(l)}(r)| \leq l!M^l$, para todo entero $l \geq l_0$ con $r > 0$ y M una constante positiva, entonces existe una*

constante $\delta > 0$ tal que

$$\|f(x) - s(x)\|_{L^\infty(\Omega)} \leq e^{\frac{-\delta}{h_{X,\Omega}}} \|f\|_{\mathcal{N}_{\Phi}(\Omega)}, \quad (41)$$

para toda $f \in \mathcal{N}_{\Phi}(\Omega)$ y cualquier conjunto de nodos X con distancia de llenado $h_{X,\Omega}$ lo suficientemente pequeña. Más aún, si ψ satisface $|\psi^{(l)}(r)| \leq M^l$, entonces

$$\|f(x) - s(x)\|_{L^\infty(\Omega)} \leq e^{\frac{-\delta |\log h_{X,\Omega}|}{h_{X,\Omega}}} \|f\|_{\mathcal{N}_{\Phi}(\Omega)}, \quad (42)$$

siempre que $h_{X,\Omega}$ sea suficientemente pequeña.

Ejemplo 1.49. Sea $\Phi(x) = e^{-\epsilon \|x\|^2}$, $\epsilon > 0$, se tiene $\psi(r) = e^{-\epsilon r}$, por lo tanto $\psi^{(l)}(r) = (-1)^l \epsilon^l e^{-\epsilon r}$ para $l \geq l_0 = 0$. Tomando $M = \epsilon$ se cumple el segundo supuesto del Teorema 1.48, es decir, para los Gaussianos tenemos límites de error de la forma (42).

Ejemplo 1.50. La función multicuadrática $\Phi(x) = (c^2 + \|x\|^2)^\beta$, $c, \beta > 0$, $\beta \notin \mathbb{N}$, se tiene $\psi(r) = (c^2 + r)^\beta$, por lo tanto

$$\psi^{(l)}(r) = \beta(\beta - 1) \cdots (\beta - l + 1)(c^2 + r)^{\beta-l}. \quad (43)$$

Luego, para $j \in \mathbb{N}$

$$\left| \frac{\beta - j}{j + 1} \right| = \left| \frac{j + 1 - (\beta + 1)}{j + 1} \right| \leq 1 + \frac{|\beta + 1|}{j + 1} \leq 1 + |\beta + 1| := M \quad (44)$$

de las ecuaciones (43) y (44) se tiene lo siguiente:

$$\begin{aligned} |\psi^{(l)}(r)| &= |\beta(\beta - 1) \cdots (\beta - l + 1)(c^2 + r)^{\beta-l}| \\ &\leq |\beta(\beta - 1) \cdots (\beta - l + 1)| \\ &\leq l! \frac{|\beta|}{1} \frac{|\beta - 1|}{2} \cdots \frac{|\beta - l + 1|}{l} \leq l! M^l, \end{aligned}$$

para $j = 0, 1, \dots, l-1$, siempre que $l \leq \lceil \beta \rceil$. Por lo tanto, las funciones multicuadráticas satisfacen la primera suposición sobre ψ , que conduce a límites de error de la forma (41).

La determinación de este orden de convergencia fue realizado por Madych y Nelson¹³, el cual puede expresarse como $O(\eta^{1/h})$ con $0 < \eta < 1$. En la Tabla 2 el valor de δ es positivo y es independiente del valor de la distancia de llenado h .

Tabla 2. Estimaciones de error para varios interpolantes FBR en términos de distancia de llenado h .

Nombre	$\phi(r)$	Estimación de error
Función multicuadrática Inversa	$(c^2 + r^2)^\beta, c > 0, \beta < 0$	$e^{-\frac{\delta}{h}}$
Función multicuadrática	$(-1)^{[\beta]}(c^2 + r^2)^\beta, c \geq 0, \beta > 0$	$e^{-\frac{\delta}{h}}$
Gaussiana	$e^{-\epsilon^2 r^2}, \epsilon > 0$	$e^{-\frac{\delta \ln(h)}{h}}$

De lo visto anteriormente, las funciones multicuadrática y gaussiana tienen un orden de convergencia exponencial. Se concluye que con base al núcleo multicuadrático se obtiene la mayor exactitud numérica, aunque la función gaussiana comparte el orden de convergencia de la función multicuadrática, es común usar el núcleo multicuadrático ya que su parámetro de forma es menos sensible¹⁴.

Ninguno de los límites de error discutidos hasta ahora ha tenido en cuenta la posibilidad de variar el parámetro de forma c para un conjunto de datos fijo X . Sin embargo, en la literatura, las funciones básicas de clase C^∞ , como la función gaussiana y la función multicuadrática (inversa), generalmente se formulan incluyendo el parámetro de forma c y uno puede preguntarse cómo un cambio en este parámetro de forma afecta las propiedades de convergencia del interpolante FBR. Se estima que el parámetro de forma depende principalmente de la función de distribución de nodos, y en consecuencia de la geometría del dominio¹⁵. Sin embargo, al día de hoy la determinación

¹³ W. R. Madych y S. A. Nelson. "Bounds on multivariate polynomials and exponential error estimates for multiquadric interpolation," en: *Journal of Approximation Theory*, 70(1) (1992), págs. 94-114.

¹⁴ R. Franke. "Scattered data interpolation: test of some methods," en: *In Mathematics of Computation*, 38 (1982), págs. 181-200.

¹⁵ H.D. Cheng. "Multiquadric and its shape parameter-A numerical investigation of error estimate, condition number, and round-off error by arbitrary precision computation," en: *Engineering Analysis with Boundary Elements* 36 (2012), págs. 220-239.

exacta y estimaciones de este parámetro en términos del número de nodos y su distribución son un problema abierto. Madych ¹⁶ estimó una cota de error que incorpora el parámetro de forma c para el núcleo multicuadrático

$$O(e^{\delta c} \eta^{\frac{c}{h}}), \quad (45)$$

con $\delta > 0$ y $0 < \eta < 1$. La convergencia de los métodos FBR se puede discutir en términos de dos tipos diferentes de aproximación: estacionaria y no estacionaria. En aproximación estacionaria, el número de centros N es fijo y el parámetro de forma c se incrementa. Este tipo de convergencia es exclusivo de los métodos FBR, ya que no existe un paralelo en los métodos basados en polinomios. La aproximación no estacionaria fija el valor del parámetro de forma y N se incrementa, como se haría al examinar la convergencia de métodos basados en polinomios.

Para un parámetro de forma fijo c , el interpolante MQ converge a una función subyacente suficientemente suave a una tasa espectral a medida que disminuye la distancia de llenado (que normalmente corresponde a un aumento en N). Es decir, el error se comporta como en (41), donde la estimación de error está sujeta a una constante que depende del valor del parámetro de forma. La estimación (41) es útil para la interpolación no estacionaria, sin embargo, no especifica cómo varía la constante δ con respecto a c , lo que disminuye su utilidad para cuantificar la convergencia de la interpolación estacionaria. Por el contrario, la estimación (45) es útil tanto para la interpolación estacionaria como para la no estacionaria, ya que muestra que la convergencia espectral resulta cuando el número de centros o parámetro de forma van a infinito.

La diferencia entre reducir h y aumentar c radica en que al incrementar el número de nodos de manera uniforme se obtiene un sistema de ecuaciones cada vez más grande y el costo computacional se incrementa considerablemente, y al aumentar c

¹⁶ W. R. Madych. "Miscellaneous error bounds for multiquadric and related interpolato," en: *Computers and Mathematics with Applications*, 24(12) (1992), págs. 121-138.

se obtiene un mal condicionamiento del sistema lineal de ecuaciones (3). Determinar el parámetro de forma óptimo ha sido un reto en los últimos años. En la siguiente sección enunciamos brevemente algunos resultados sobre la determinación de dicho parámetro.

1.6. PARÁMETRO DE FORMA

La precisión de los métodos sin malla basados en FBR dependen en gran medida de los centros de base radial y del llamado parámetro de forma c . Las FBR comúnmente usadas son las funciones multicuadráticas y la función gaussiana, debido a su rápida tasa de convergencia. Sin embargo, la precisión de los métodos sin malla basados en las funciones anteriormente mencionadas está muy influenciada por el parámetro de forma. A medida que c se hace grande, estas FBR se vuelven planas y la matriz resultante se vuelve mal condicionada ¹⁵. La precisión de las FBR continúa mejorando a medida que c aumenta; sin embargo, cuando c es lo suficientemente grande, la precisión empeora y finalmente se rompe. Este fenómeno ha sido descrito por Schaback ¹⁷ como un “principio de incertidumbre”.

En la literatura se han propuesto distintos enfoques para seleccionar c , principalmente para el núcleo multicuadrático y datos en dos dimensiones. Nuestro interés está enfocado en el método de colocación asimétrico basado en el núcleo multicuadrático. En los últimos años, se han realizado varios intentos para encontrar, mediante experimentos numéricos o argumentos heurísticos, un parámetro de forma “óptimo”. La estrategia más usada, es evaluar el interpolante MQ en un rango del parámetro de forma y llamar óptimo al valor que da como resultado el error más pequeño. Además de ser ineficiente, la estrategia solo funciona si se conoce la función que se está aproximando. Las primeras fórmulas para valores de parámetros de forma “bueno” dependían

¹⁷ R. Schaback. “Error estimates and condition numbers for radial basis function interpolation,” en: *Advances in Computational Mathematics*, 3 (1995), págs. 251-264.

del espacio entre centros. En un trabajo temprano, Hardy ¹⁸ recomendó usar

$$c = 0.815d \text{ donde } d = \frac{1}{N} \sum_{i=1}^N d_i,$$

y d_i es la distancia desde el i -ésimo centro a su vecino más cercano.

Más tarde, Franke ¹⁴ ofreció una fórmula similar

$$c = (1.25D)/\sqrt{N},$$

donde D es el diámetro del círculo más pequeño que encierra todos los centros.

Se han realizado varios intentos para determinar valores buenos del parámetro de forma, como en las referencias ¹⁹ y ²⁰. Muchos de los intentos de proporcionar valores buenos del parámetro de forma reflejan el principio de incertidumbre. La estrategia comúnmente usada desarrollada mediante experimentos numéricos es producir una matriz de sistema con un número de condición suficientemente grande para proporcionar una buena precisión, pero no demasiado alto para que los errores comiencen a dominar.

Heurísticamente se ha argumentado que usar un parámetro de forma variable es una buena idea. Un parámetro de forma variable se refiere al uso de un valor diferente del parámetro de forma en cada centro. Esto da como resultado parámetros de forma que son los mismos en cada columna de la matriz de interpolación. Un argumento para usar un parámetro de forma variable es que conduce a entradas distintas en las matrices FBR, lo que a su vez conduce a números de condición más bajos. Una

¹⁸ R. L. Hardy. "Multiquadric equations of topography and other irregular surfaces," en: *Journal of Geophysical Research*, 76(8) (1971), págs. 1905-1915.

¹⁹ T. A. Foley. "Near optimal parameter selection for multiquadric interpolation," en: *Journal of Applied Science and Computation*, 1 (1994), págs. 54-69.

²⁰ C. J. Trahan y R. E. Wyatt. "Radial basis function interpolation in the quantum trajectory method: optimization of the multiquadric shape parameter," en: *Journal of Computational Physics*, 185 (2003), págs. 27-49.

consecuencia negativa de usar una forma variable es que la matriz de interpolación ya no es simétrica.

En ²¹, la fórmula

$$c_j^2 = c_{\min}^2 \left[\frac{c_{\max}^2}{c_{\min}^2} \right]^{\frac{j-1}{N-1}}, \quad j = 1, \dots, N, \quad (46)$$

da un parámetro de forma que varía exponencialmente. Otra estrategia de parámetros de forma variable es el parámetro de variación lineal

$$c_j = c_{\min} + \left(\frac{c_{\max} - c_{\min}}{N - 1} \right) j, \quad j = 0, 1, \dots, N - 1.$$

En la referencia ²¹ se pueden encontrar más experimentos numéricos con la estrategia de parámetros de forma variable. Cheng ²² con base en un problema elíptico, observa numéricamente que el error residual tomado en una muestra de los datos, se puede utilizar como un indicador de error para optimizar la selección de c .

En el apéndice A se describe el número de condición y el mal condicionamiento de una matriz.

²¹ E. J. Kansa y R. E. Carlson. "Improved accuracy of multiquadric interpolation using variable shape parameter," en: *Computers and Mathematics with Applications*, 24 (1992), págs. 99-120.

²² Golberg M. A. Kansa E. J. Cheng H.D. y G. Zammuto. "Exponential convergence and h-c multiquadric collocation method for partial differential equations," en: *Numerical Methods Partial Differential Equations*, 19(5) (2003), págs. 571-594.

2. ESQUEMA DE COLOCACIÓN ASIMÉTRICO

El presente capítulo tiene como objetivo la descripción del método de colocación asimétrico a través de su aplicación a la resolución numérica de ecuaciones en derivadas parciales estacionarias de tipo elíptico. Analizaremos numéricamente la ecuación de Poisson con condición de frontera tipo Dirichlet y Neumann. Además se analizará la tasa de convergencia del esquema ya nombrado, y se hará una comparación del parámetro de forma para las aproximaciones.

2.1. PROBLEMA LINEAL ESTACIONARIO

El esquema numérico propuesto por E. J. Kansa en 1990 ¹² busca aproximar las derivadas parciales por medio de funciones de base radial. Se menciona que el método FBR es uno de los más utilizados para la solución numérica de ecuaciones diferenciales parciales que surgen en problemas de ingeniería. Su popularidad se debe a la no utilización de malla, lo que significa que solo se requiere un conjunto de puntos para la discretización del problema continuo. Esto hace que la implementación del método sea particularmente fácil, especialmente para problemas con geometrías complejas en tres dimensiones. Una desventaja del método es la elección óptima del parámetro de forma que se encuentra en la mayoría de los FBR. Además, la discretización de los métodos de colocación FBR conduce a matrices muy mal condicionadas y esto ha limitado la precisión a un cierto nivel.

La siguiente sección expone el esquema general de la aproximación usando el método de colocación asimétrico aplicado a ecuaciones en derivadas parciales lineales no dependientes del tiempo. Con el fin de no expandirnos en el documento se enuncia el método de colocación asimétrico basado en funciones definidas positivas.

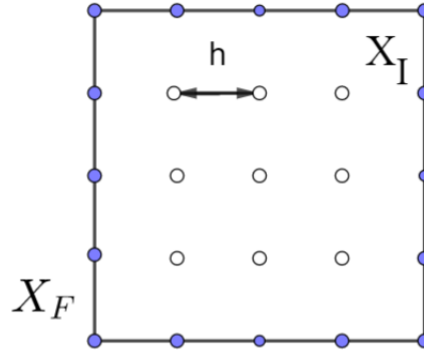
Considere el siguiente problema de valor en la frontera

$$\begin{cases} \mathcal{L}u(x) = f(x), & x \in \Omega, \\ \mathcal{B}u(x) = g(x), & x \in \partial\Omega, \end{cases} \quad (47)$$

donde Ω es un dominio acotado de $\mathbb{R}^d$ con frontera $\partial\Omega$, $\mathcal{L}$ es un operador diferencial lineal, $\mathcal{B}$ denota la aplicación de las condiciones de frontera (Dirichlet, Neumann o Robin) y $f, g : \mathbb{R}^d \rightarrow \mathbb{R}$ son funciones conocidas.

En el esquema de colocación asimétrico se considera que el dominio Ω es discretizado mediante N distintos nodos de colocación denotados por $X = X_I \cup X_F$, donde $X_I = \{x_1, x_2, \dots, x_{N_I}\}$ es el conjunto de nodos interiores y $X_F = \{x_{N_I+1}, \dots, x_N\}$ es el conjunto de nodos que conforman la frontera, note que $N_F = N - N_I$. En la Figura 4 se ilustra la distribución de los nodos.

Figura 4. Distribución nodos 2D.



Como es bien sabido, en general no es posible determinar la solución analítica de ecuaciones en derivadas parciales; es por ello la necesidad de aproximar numéricamente la solución exacta $u(x)$ de la ecuación diferencial parcial (47). Para determinar la solución aproximada $\tilde{u}(x)$ hacemos uso del interpolante radial, donde $\tilde{u}(x)$ puede ser representada como una combinación lineal de funciones radiales

$$\tilde{u}(x) = \sum_{j=1}^N \lambda_j \Phi(\|x - x_j\|), \quad x \in X. \quad (48)$$

Ahora, sustituyendo la aproximación radial $\tilde{u}(x)$ en el problema (47) obtenemos

$$\begin{aligned}\mathcal{L}\tilde{u}(x_i) &= \sum_{j=1}^N \lambda_j \mathcal{L}\Phi(\|x_i - x_j\|) = f(x_i), \quad x_i \in X_I, \\ \mathcal{B}\tilde{u}(x_i) &= \sum_{j=1}^N \lambda_j \mathcal{B}\Phi(\|x_i - x_j\|) = g(x_i), \quad x_i \in X_F.\end{aligned}$$

Las igualdades anteriores quedan expresadas en un sistema matricial definido como

$$\begin{bmatrix} \mathcal{L}\Phi \\ \mathcal{B}\Phi \end{bmatrix} \begin{bmatrix} \lambda \end{bmatrix} = \begin{bmatrix} f \\ g \end{bmatrix}, \quad (49)$$

donde el vector de incógnitas $\lambda \in \mathbb{R}^N$, $\mathcal{L}\Phi \in \mathbb{R}^{N_I \times N}$ y $\mathcal{B}\Phi \in \mathbb{R}^{N_F \times N}$. El método descrito es llamado *colocación asimétrico* debido a la asimetría de la matriz resultante en la discretización. El análisis de la no singularidad del sistema (49) es un problema que posee preguntas abiertas y se constituye en uno de los mayores retos actuales dentro de la comunidad académica de esta área de estudio. La solución del sistema anterior se puede expresar como:

$$\lambda = \begin{bmatrix} \mathcal{L}\Phi \\ \mathcal{B}\Phi \end{bmatrix}^{-1} \begin{bmatrix} f \\ g \end{bmatrix},$$

siempre que la inversa de $\begin{bmatrix} \mathcal{L}\Phi \\ \mathcal{B}\Phi \end{bmatrix}$ exista. Con base en la solución obtenida λ podemos reconstruir la solución aproximada $\tilde{u}(x)$ a través de (48).

En el estudio del método de colocación asimétrico, Kansa propuso específicamente el uso de funciones multicuadráticas en (48) y, en consecuencia, este enfoque de colocación asimétrico a menudo aparece en la literatura como el método multicuadrático (o método Kansa-MQ). En las siguientes secciones se describe con algo de detalle el paso a paso de la implementación del método asimétrico para el problema lineal estacionario usando como núcleo el multicuadrático.

2.2. Problema lineal estacionario

El objetivo de esta sección es evidenciar mediante un problema lineal estacionario la aplicación del método Kansa-MQ. En primer lugar, describimos el método aplicado a un problema de Poisson con condiciones de frontera Dirichlet y Neumann en un dominio rectangular. En segundo lugar, se evidencia la eficiencia del método asimétrico basado en el error cuadrático medio y error máximo entre la solución analítica y la solución numérica para un número de centros N y parámetro de forma c fijos. Finalmente, en aras de mejorar la eficiencia del método, se analizan los errores ya mencionados para valores diferentes de c y N cuando uno de los dos está fijo.

Consideremos el siguiente problema de Poisson en el dominio rectangular $\Omega = (0, 1) \times (0, 1)$,

$$\frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2} = f(x, y), \quad (x, y) \in (0, 1) \times (0, 1), \quad (50)$$

donde $f(x, y) = -2(2y^3 - 3y^2 + 1) + 6(2y - 1)(1 - x^2)$, sujeta a condiciones de frontera tipo Dirichlet

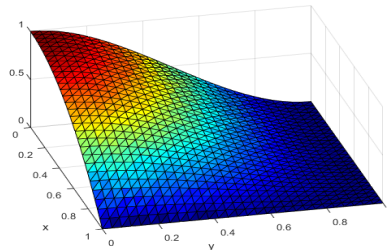
$$u(0, y) = 2y^3 - 3y^2 + 1, \quad u(1, y) = 0, \quad 0 \leq y \leq 1, \quad (51)$$

y tipo Neumann

$$\frac{\partial u}{\partial y} = 0, \quad \text{a lo largo de } y = 0, y = 1. \quad (52)$$

Figura 5. Solución analítica del problema de Poisson (50)

$$u(x, y) = (1 - x^2)(2y^3 - 3y^2 + 1).$$



Una forma de medir la calidad de las aproximaciones calculadas, es utilizando el error cuadrático medio (ECM), expresado como:

$$ECM = \sqrt{\frac{1}{N} \sum_{i=1}^N (u_i - \tilde{u}_i)^2}, \quad (53)$$

y el error máximo

$$E_{\text{máx}} = \|u - \tilde{u}\|_{L^\infty}. \quad (54)$$

En base a los dos errores descritos anteriormente, discutiremos la eficiencia del método asimétrico. Consideraremos una distribución inicial uniforme X con $N = 100$ nodos, divididos en 64 nodos en el interior y 36 nodos en la frontera; además, la función radial a usar es la función multicuadrática para un parámetro de forma $c = 2$, es decir,

$$\phi(\|(x - x_j, y - y_j)\|) = \sqrt{(x - x_j)^2 + (y - y_j)^2 + c^2}, \quad (x_j, y_j) \in X. \quad (55)$$

De la sección anterior, recordamos que la solución aproximada $\tilde{u}$ se define como

$$\tilde{u}(x, y) = \sum_{j=1}^N \alpha_j \phi(\|(x - x_j, y - y_j)\|). \quad (56)$$

Las derivadas parciales de $\tilde{u}$ en (x_i, y_i) están dadas por

$$\begin{aligned} \frac{\partial \tilde{u}}{\partial x}(x_i, y_i) &= \sum_{j=1}^N \alpha_j \left[\frac{(x_i - x_j)}{\phi_{ij}} \right], \\ \frac{\partial \tilde{u}}{\partial y}(x_i, y_i) &= \sum_{j=1}^N \alpha_j \left[\frac{(y_i - y_j)}{\phi_{ij}} \right], \\ \frac{\partial^2 \tilde{u}}{\partial x^2}(x_i, y_i) &= \sum_{j=1}^N \alpha_j \left[\frac{1}{\phi_{ij}} - \frac{(x_i - x_j)^2}{\phi_{ij}^3} \right], \\ \frac{\partial^2 \tilde{u}}{\partial y^2}(x_i, y_i) &= \sum_{j=1}^N \alpha_j \left[\frac{1}{\phi_{ij}} - \frac{(y_i - y_j)^2}{\phi_{ij}^3} \right], \end{aligned}$$

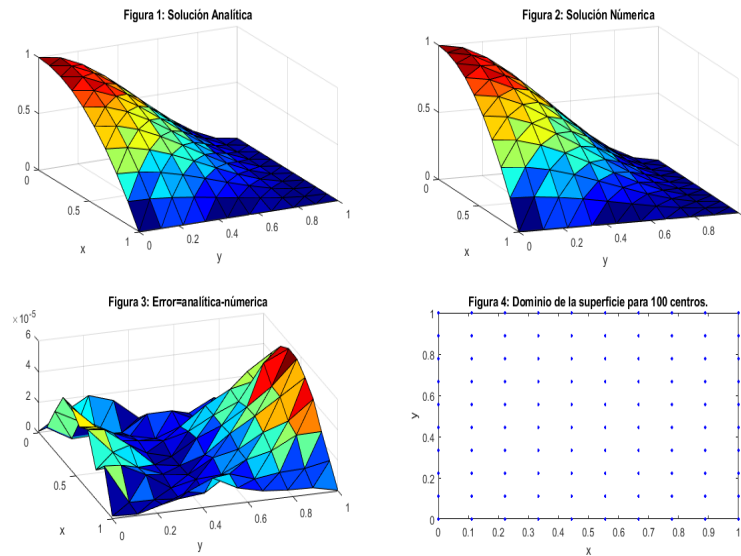
donde $\phi_{ij} := \phi(x_i - x_j, y_i - y_j)$. Sustituyendo estas derivadas parciales en (50), encon-

tramos, para cada $(x_i, y_i), (x_j, y_j) \in X$,

$$\begin{aligned}
\sum_{j=1}^N \alpha_j \left[\frac{2}{\phi_{ij}} - \frac{((x_i - x_j)^2 + (y_i - y_j)^2)}{\phi_{ij}^3} \right] &= f(x_i, y_i), \\
\sum_{j=1}^N \alpha_j \phi(0 - x_j, y_i - y_j) &= 2y_i^3 - 3y_i^2 + 1, \\
\sum_{j=1}^N \alpha_j \phi(1 - x_j, y_i - y_j) &= 0, \\
\sum_{j=1}^N \alpha_j \left[\frac{(0 - y_j)}{\phi_{ij}} \right] &= 0, \\
\sum_{j=1}^N \alpha_j \left[\frac{(1 - y_j)}{\phi_{ij}} \right] &= 0.
\end{aligned} \tag{57}$$

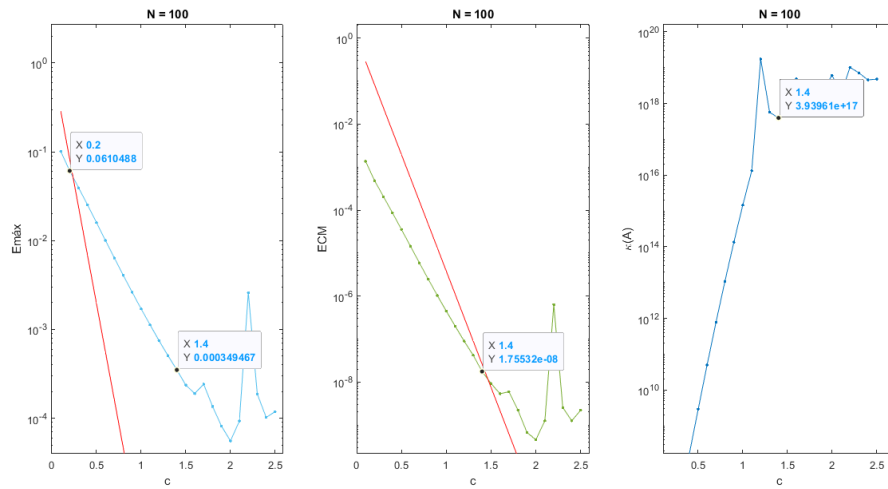
Resolviendo el sistema (57) obtenemos una solución aproximada del problema diferencial. La Figura 6 muestra la solución aproximada del problema de Poisson usando el método asimétrico para los parámetros c y N definidos previamente, con errores $ECM = 4,5337 \times 10^{-10}$ y $E_{\text{máx}} = 5,5371 \times 10^{-5}$.

Figura 6. Figura 1: Representa la solución analítica del problema de Poisson (50). Figura 2: Representa la solución numérica del problema de Poisson (50) usando el método asimétrico. Figura 3: Expone el error puntual entre la solución analítica y solución numérica. Figura 4: Describe el dominio para 100 centros distribuidos uniformemente.



En aras de reducir el error en la aproximación implementando el método de colocación asimétrico se analizarán los dos tipos de aproximación detallados en la Sección 1.5. Primero analizaremos la aproximación estacionaria fijando el número de nodos $N = 100$ e incrementando el parámetro de forma c uniformemente a razón 0.1. Los resultados de la aproximación estacionaria se observan en la Figura 7.

Figura 7. Método de aproximación estacionaria: N se fija en $N = 100$. Izquierda: Error máximo frente al parámetro de forma del problema de Poisson (azul). Límite de error de la forma 45 (rojo). Centro: Error cuadrático medio contra el parámetro de forma (verde). Límite de error de la forma 45 (rojo). Derecha: Número de condición del sistema (57) vs el parámetro de forma.



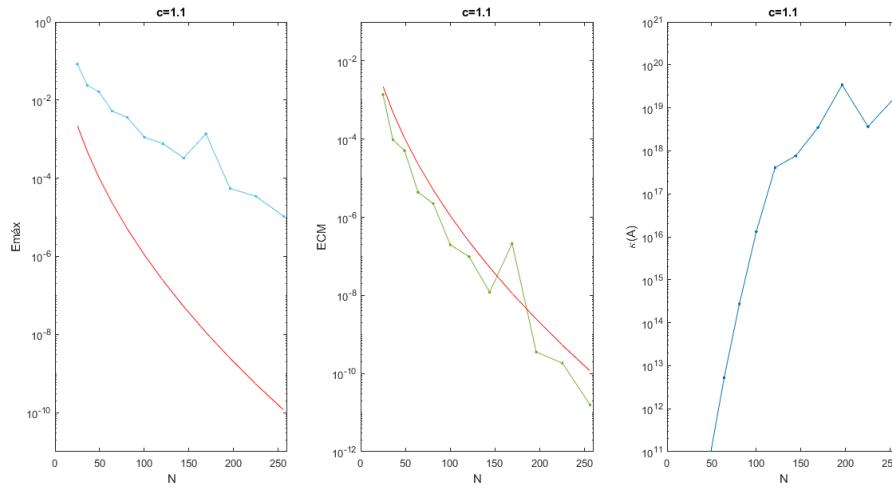
Las dos primeras gráficas de la Figura 7 describen los errores $E_{\max}$ y ECM respectivamente; además se traza un límite de error de la forma $\eta^{c/h}$ para $\eta = 1/4$, ver (45). La estimación del error se mantiene y la tasa de convergencia esperada se logra para valores de $c \leq 0.2$ cuando se mide la eficacia del método basado en el error máximo ($E_{\max}$). Para el error cuadrático medio (ECM) la estimación del error se mantiene y la tasa de convergencia esperada se logra para valores de $c \leq 1.4$.

En la aproximación estacionaria se lograron los errores mínimos para un parámetro de forma $c = 2$. Sin embargo, $c = 2$ no obedece la tasa convergencia esperada. Por tal razón, se verifica el principio de incertidumbre de Schaback, en efecto para $c = 2$ se obtuvo un error mínimo, pero se generó un mal condicionamiento del sistema lineal

(57).

A continuación se analiza la aproximación no estacionaria, seleccionando a c de tal forma que verifique la tasa de convergencia, ver Figura 8.

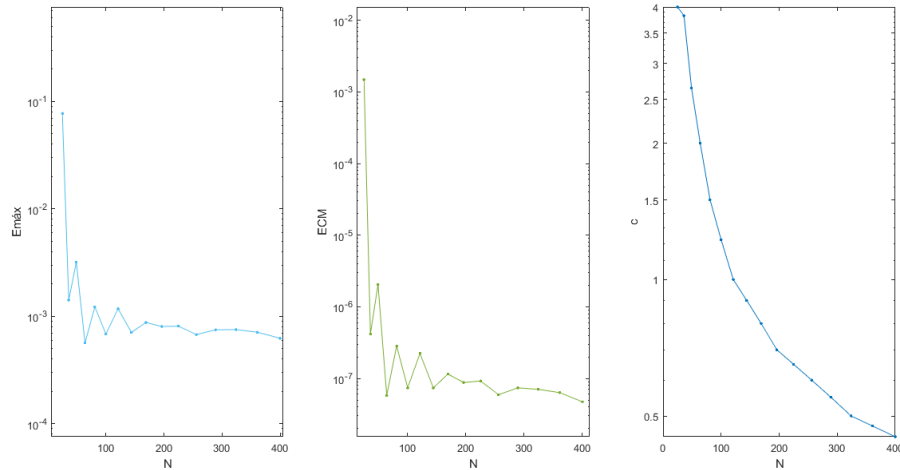
Figura 8. Método de aproximación no estacionaria: El parámetro de forma se fija en $c = 1.1$. Izquierda: Error máximo frente al número de centros (azul). Límite de error de la forma 45 (rojo). Centro: Error cuadrático medio contra el número de centros N (verde). Límite de error de la forma 45 (rojo). Derecha: Número de condición del sistema (57) vs número de centros



La aproximación no estacionaria al igual que la aproximación estacionaria muestran que un aumento de N o c generan un mal condicionamiento del sistema lineal (57). Sin embargo, para la aproximación no estacionaria el mínimo error calculado por ECM se obtuvo para N tal que verifica la tasa de convergencia esperada.

Una pregunta que puede surgir a partir de la aproximación estacionaria y no estacionaria es ¿Cómo se comporta la estimación de error cuando N y c varían al tiempo? En realidad, este es el enfoque que se adopta con más frecuencia en las aplicaciones. Solucionemos de nuevo el problema de Poisson (50) para $N = 9, 16, 25, 36, \dots, 400$ y ajustemos el parámetro de forma de tal manera que el número de condición del sistema (57) esté entre 1×10^{15} y 1×10^{17} . Lo anterior basado en los resultados obtenidos de la aproximación estacionaria y no estacionaria.

Figura 9. Interpolación del problema de Poisson variando tanto N como c . Izquierda: Error máximo frente al número de centros. Centro: Error cuadrático medio contra el número de centros N . Derecha: Parámetro de forma vs número de centros.



Los resultados se ilustran en la Figura 9, se observa una tasa de convergencia exponencial de $N = 25$ a $N = 400$. Los errores mínimos se dan para $N = 400$. Note que a medida que N se hace grande, el parámetro de forma se debe reducir para que el número de condición del sistema sea deseable. Si bien el parámetro de forma y el número de centros particularmente dependen del problema, este resultado es típico. Dado que los centros están espaciados uniformemente, el parámetro de forma disminuye a una tasa exponencial simple con N para mantener el acondicionamiento deseado. Con centros dispersos, no existiría una relación tan simple.

2.3. ELEMENTOS FINITOS VS MÉTODO DE COLOCACIÓN ASIMÉTRICO

Uno de los objetivos que tiene esta monografía es intentar verificar la eficiencia de los métodos libres de malla en FBR sobre los métodos tradicionales, como es el caso de elementos finitos. En esta sección resolvemos el problema de Poisson (50)-(52) mediante el método de elementos finitos para poder compararlo con el método de colocación asimétrico usando la función radial multicuadrática.

El resultado de solucionar el problema de Poisson mediante el método de elementos finitos se puede visualizar en la Figura 11. La implementación numérica es realizada

mediante la función *pdetool* de Matlab la cual nos da por defecto un mallado de $N = 1541$ nodos.

Figura 10. Dominio rectangular de la superficie con 1541 nodos.

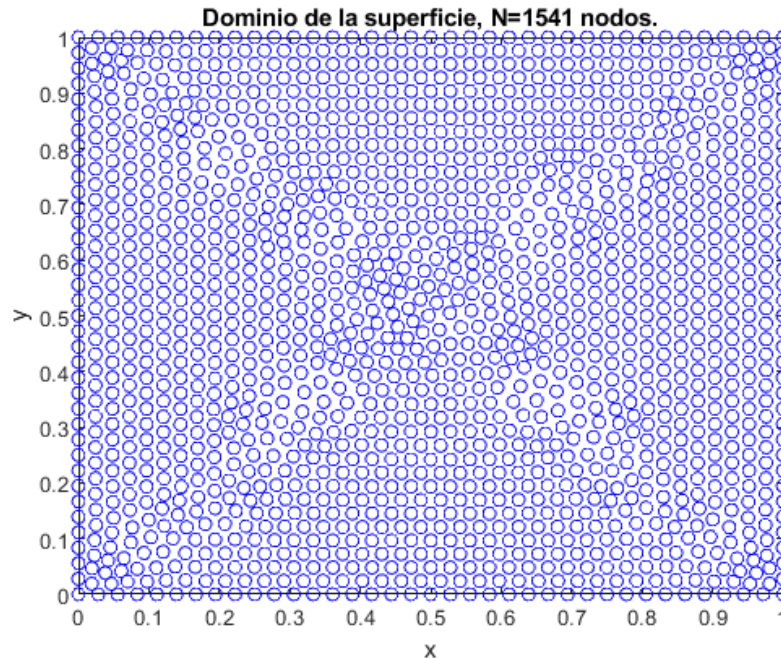
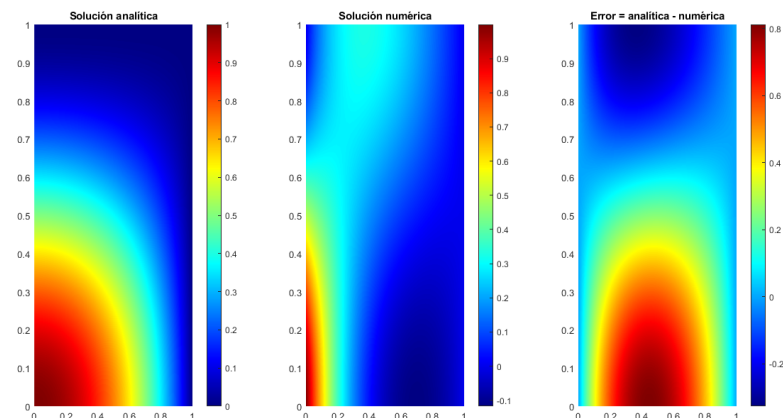


Figura 11. Solución numérica método elementos finitos. Izquierda: Solución analítica problema de Poisson (50)-(52). Centro: Solución numérica del problema de Poisson (50)-(52). Derecha: Error puntual entre la solución analítica y solución numérica del problema (50)-(52).



Para solucionar el problema (50)-(52) mediante el método de colocación asimétrico, se escogieron como centros los 1541 nodos de la malla creada para elementos finitos.

Estos están divididos en 1397 nodos interiores y 144 nodos de frontera, inicialmente consideramos un parámetro de forma $c = 0.1$. Los errores alcanzados usando los métodos ya mencionados se pueden apreciar en la Tabla 3.

Tabla 3. Comparación método de colocación asimétrico vs método elementos finitos.

Método	N	ECM	$E_{\text{máx}}$
Método colocación asimétrico	1541	1.2359×10^{-6}	0.0036
Método elementos finitos	1541	0.1140	0.8123

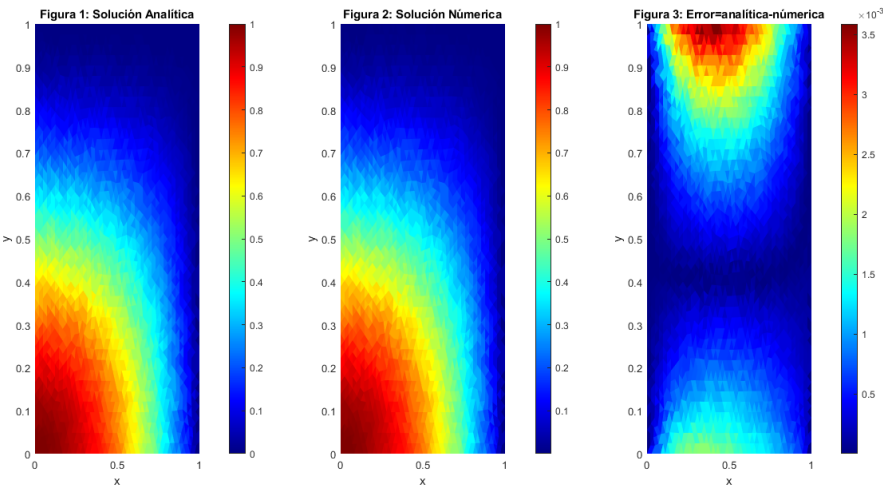
Note que el método de colocación asimétrico tiene orden superior al método de elementos finitos, esto es consistente con la convergencia espectral reportada para problemas elípticos. Además se puede observar que conforme crece el valor de c , el error disminuye, ver Tabla 4.

Tabla 4. Comparación método colocación asimétrico vs método elementos finitos.

Método	c	ECM	$E_{\text{máx}}$
Método	0.2	7.1279×10^{-9}	2.7019×10^{-4}
Colocación asimétrico	0.3	6.3519×10^{-11}	2.68×10^{-5}

La solución numérica del problema de Poisson usando el método colocación asimétrico con parámetro de forma $c = 0.1$ se ilustra en la Figura 12.

Figura 12. Solución numérica método de colocación asimétrico con parámetro de forma $c = 0.1$. La imagen se divide en solución analítica (izquierda), solución numérica (centro) del problema de Poisson descrito en (50) y error puntual entre las soluciones antes mencionadas (derecha).



Terminamos esta sección con las siguientes observaciones a manera de conclusiones. En el método de colocación asimétrico (o MQ-Kansa) se pueden observar posibles dificultades de implementación numérica; en primer lugar, cuando el número de centros N es grande, la dificultad numérica es la eficiencia. En este caso, este problema se puede solucionar con métodos de descomposición de dominios ²³ y con métodos iterativos que utilizan técnicas de preacondicionamiento ^{24,25}. Además, se puede utilizar un algoritmo más codicioso para reducir el número de centros ²⁶. Por otro lado, cuando se tiene un parámetro de forma grande, la dificultad numérica surge por problemas de acondicionamiento. Este se puede abordar con precisión numérica extendida ²⁷, o evaluando el método MQ-Kansa de una manera que pase por alto los sistemas lineales mal acondicionados (Algoritmo Contour-Padé)^{28,29}, o usando algoritmos distintos a la eliminación gaussiana para resolver los sistemas lineales (La descomposición del valor singular) ³⁰. Cabe aclarar que estas dificultades se pueden dar en conjunto o por separado.

²³ J. Li e Y. C. Hon. "Domain decomposition for radial basis meshless methods," en: *Numerical Methods for Partial Differential Equations*, 20(3) (2004), págs. 450-462.

²⁴ L. Ling y E.J. Kansa. "Preconditioning for Radial Basis Functions with Domain Decomposition Methods," en: *Mathematical and Computer Modelling, Submitted to Mathematical and Computer Modeling*, 40 (2004), págs. 1413-1427.

²⁵ L. Ling y Kansa. E.J. "A least squares preconditioner for radial basis functions collocation methods," en: *Advances in Computational Mathematics*, 23 (2005), págs. 31-54.

²⁶ Schaback. R De Marchi. S y Wendland. H. "Near-optimal data-independent point locations for radial basis function interpolation," en: *Advances in Computational Mathematics*, 23 (2005), págs. 317-330.

²⁷ Leeb C.F. Huang C.S y A.H.D. Cheng. "Error estimate, optimal shape factor, and high precision computation of multiquadric collocation method," en: *Engineering Analysis with Boundary Elements*, 31 (2007), págs. 614-623.

²⁸ G. Wright. *Radial Basis Function Interpolation: Numerical and Analytical Developments, PhD thesis*,

²⁹ B. Fornberg y G. Wright. "Stable computation of multiquadric interpolants for all values of the shape parameter," en: *Computers and Mathematics with applications*, 48 (2004), págs. 853-867.

³⁰ L. N. Trefethen y D. Bau. *Numerical Linear Algebra, first edition*,

3. MÉTODO DE COLOCACIÓN ASIMÉTRICO APLICADO A LA ECUACIÓN DEL CALOR

El interés del presente capítulo es comprobar la eficiencia del método asimétrico usando la función radial multicuadrática para ecuaciones lineales dependientes del tiempo. Primero se expone la generalización del método asimétrico a un problema evolutivo, junto con el análisis de estabilidad del problema dependiente del tiempo. Finalmente, mediante el método de colocación asimétrico solucionamos un problema parabólico en dos dimensiones. Este enfoque de colocación asimétrico a menudo aparece en la literatura como el método multicuadrático (o método Kansa-MQ).

3.1. GENERALIDADES

El siguiente capítulo consiste en describir el método de colocación asimétrico, aplicado a las ecuaciones diferenciales lineales dependientes del tiempo, el cual se vuelve más complejo debido a que se debe considerar la variación temporal. Con el fin de esclarecer detalles se ejemplifica el método usando la ecuación del calor lineal en dos dimensiones.

Considere el siguiente sistema de ecuaciones, cuya ecuación diferencial es una ecuación parabólica general dependiente del tiempo, sujeta a las condiciones de frontera y condición inicial respectivamente

$$u_t(x, t) - \mathcal{L}u(x, t) = f(x, t), \quad x \in \Omega, t \in (0, T), \quad (58)$$

$$\mathcal{B}u(x, t) = g(x, t), \quad x \in \partial\Omega, t \in (0, T), \quad (59)$$

$$u(x, 0) = u_0(x), \quad x \in \Omega, \quad (60)$$

donde Ω es un subconjunto abierto y acotado de $\mathbb{R}^d$, y f y g son funciones dadas, u_0 corresponde al dato inicial, $\mathcal{L}$ es un operador diferencial parcial de tipo elíptico, y $\mathcal{B}$ es un operador de frontera que puede ser de tipo Dirichlet o Neumann.

Usando la notación antes vista, la solución aproximada $\tilde{u}(x, t)$ de la solución exacta

$u(x, t)$ del problema de valor inicial (58)-(60) es definida como una combinación lineal de funciones radiales, de la siguiente forma:

$$\tilde{u}(x, t) = \sum_{j=1}^N \lambda_j(t) \phi(\|x - x_j\|), \quad x, x_j \in X, \quad (61)$$

donde los $\lambda_j(t)$ son los coeficientes dependientes del tiempo a ser determinados, y ϕ denota la función de base radial. El problema (58)-(60), a diferencia del problema estacionario descrito en el capítulo anterior, consiste en la adición de la variable temporal en la ecuación diferencial parcial y por tanto en la aproximación $\tilde{u}$.

La discretización espacial del problema evolutivo (58)-(60) usando el método asimétrico está definida como:

$$\mathcal{L}\tilde{u}(x) = \sum_{j=1}^N \lambda_j(t) \mathcal{L}\phi(\|x - x_j\|), \quad x, x_j \in X, \quad (62)$$

sobre los N distintos puntos de colocación. Para la aproximación de la derivada temporal del problema lineal (58)-(60) emplearemos el “método clásico” θ , con $0 \leq \theta \leq 1$, en el sentido que, dependiendo del valor de θ en (66) se obtiene un método explícito ($\theta = 0$) o un método implícito ($\theta = 1$), o el esquema de Crank-Nicholson ($\theta = 1/2$). La ecuación resultante se puede expresar como:

$$\begin{aligned} \frac{u(x, t + \Delta t) - u(x, t)}{\Delta t} - [\theta \mathcal{L}u(x, t + \Delta t) + (1 - \theta) \mathcal{L}u(x, t)] &= f(x, t + \Delta t), \\ \mathcal{B}u(x, t + \Delta t) &= g(x, t + \Delta t), \end{aligned}$$

donde Δt denota al tamaño del paso temporal. Empleando la notación $u(x, t^n) = u^n$ y $t^n = t^{n-1} + \Delta t$, las dos ecuaciones anteriores se pueden representar como

$$\begin{aligned} u^{n+1} - \Delta t \theta \mathcal{L}u^{n+1} &= u^n + \Delta t f^{n+1} + \Delta t (1 - \theta) \mathcal{L}u^n, \\ \mathcal{B}u^{n+1} &= g^{n+1}. \end{aligned} \quad (63)$$

Para obtener la solución aproximada $\tilde{u}(x, t)$ de la solución exacta $u(x, t)$ del problema de valor inicial, empleamos la combinación lineal de funciones radiales (61) en (63)

aplicando las condiciones de frontera. Sea X el conjunto de N nodos de colocación distribuidos en nodos en el interior X_I y nodos que conforman la frontera X_F , para $x, x_j \in X$ el sistema (63) queda expresado como:

$$\begin{aligned} \sum_{j=1}^N \lambda_j^{n+1} \left(\Phi(\|x - x_j\|) - \Delta t \theta \mathcal{L} \Phi(\|x - x_j\|) \right) &= \Delta t f^{n+1} \\ + \sum_{j=1}^N \lambda_j^n \left(\Phi(\|x - x_j\|) + (1 - \theta) \Delta t \mathcal{L} \Phi(\|x - x_j\|) \right), & \\ \sum_{j=1}^N \lambda_j^{n+1} \mathcal{B} \Phi(\|x - x_j\|) &= g^{n+1}, \end{aligned} \quad (64)$$

Simplificando el sistema (64) se obtiene

$$\begin{bmatrix} \Phi_a - \Delta t \theta \mathcal{L} \Phi_a \\ \mathcal{B} \Phi_b \end{bmatrix} \lambda^{n+1} = \begin{bmatrix} \Phi_a + (1 - \theta) \Delta t \mathcal{L} \Phi_a \\ 0_b \end{bmatrix} \lambda^n + \begin{bmatrix} \Delta t f^{n+1} \\ g^{n+1} \end{bmatrix}, \quad (65)$$

donde $\lambda \in \mathbb{R}^N$ y $0_b \in \mathbb{R}^{N_F \times N}$ representa la matriz de ceros. La aplicación del operador diferencial $\mathcal{L}$ sobre el conjunto de nodos interiores es la matriz $\mathcal{L} \Phi_a \in \mathbb{R}^{N_I \times N}$ donde la matriz Φ_a representa una parte de la matriz $A = [\Phi_a \ \Phi_b] \in \mathbb{R}^{N \times N}$. La aplicación del operador diferencial $\mathcal{B}$ sobre los nodos que conforman la frontera está definida por $\mathcal{B} \Phi_b \in \mathbb{R}^{N_F \times N}$.

La siguiente expresión describe la forma compacta de (65)

$$H_+ \lambda^{n+1} = H_- \lambda^n + F^{n+1} \quad (66)$$

donde

$$H_+ = \begin{bmatrix} \Phi_a - \Delta t \theta \mathcal{L} \Phi_a \\ \mathcal{B} \Phi_b \end{bmatrix}, \quad H_- = \begin{bmatrix} \Phi_a + (1 - \theta) \Delta t \mathcal{L} \Phi_a \\ 0_b \end{bmatrix}$$

y

$$F^{n+1} = \begin{bmatrix} \Delta t f^{n+1} \\ g^{n+1} \end{bmatrix}.$$

El lado derecho en (66) representa la solución conocida al tiempo t_n , mientras que el lado izquierdo representa la solución desconocida al tiempo t_{n+1} . La ecuación (66) representa el sistema iterativo a resolver en cada paso del tiempo desde $n = 0$ hasta $n = n_{\text{máx}}$ donde $n_{\text{máx}} = t_{\text{máx}}/\Delta t$. Este sistema iterativo corresponde a la aproximación numérica del problema de valor inicial empleando el método de colocación asimétrico.

Dependiendo del valor de θ en (66) se obtiene un método explícito ($\theta = 0$), un método implícito ($\theta = 1$), o el esquema de Crank-Nicholson ($\theta = 1/2$). En la siguiente sección, se realiza un análisis de estabilidad para los métodos explícito e implícito descritos anteriormente.

3.2. ANÁLISIS DE ESTABILIDAD

A continuación se describe un criterio de estabilidad matricial para conocer el tamaño de Δt del sistema iterativo (66) tomado como referencia ^{31,32}. La solución aproximada discreta de (61) es representada como

$$u^n = A\lambda^n \quad (67)$$

siendo $A = [\Phi_a \ \Phi_b]^T$ la matriz de interpolación, λ el vector de incógnitas. Ahora despejando λ^{n+1} de (66) y haciendo $F = 0$ obtenemos

$$\lambda^{n+1} = H_+^{-1} H_- \lambda^n. \quad (68)$$

³¹ Djidjeli K. Chinchapatnam P. P y P.B. Nair. "Domain decomposition for time-dependent problems using radial based meshless methods." En: *Numerical Methods for Partial Differential Equations*, 23(1) (2007), págs. 38-59.

³² Djidjeli. K Zerroukat M y Charafi. "Explicit and implicit meshless methods for linear advection-diffusion-type partial differential equations," en: *International Journal of Numerical Methods in Engineering*, 48(1) (2000), págs. 19-35.

Sustituyendo u^n definida en (67) en la anterior ecuación se tiene que

$$u^{n+1} = Ku^n, \quad (69)$$

donde $K = AH_+^{-1}H_-A^{-1}$. Tomamos una perturbación ϵ^n en el tiempo n -ésimo, es decir, $\epsilon^n = u^n - \hat{u}^n$, donde $\hat{u}^n$ es la solución computada tal que $\epsilon^{n+1} = K\epsilon^n$. En efecto, note que

$$\begin{aligned} \epsilon^{n+1} &= u^{n+1} - \hat{u}^{n+1} \\ &= Ku^n - K\hat{u}^n \\ &= K(u^n - \hat{u}^n) \\ &= K\epsilon^n. \end{aligned} \quad (70)$$

El sistema es estable si $\epsilon^n \rightarrow 0$ conforme $n \rightarrow \infty$, siempre que el radio espectral de K sea menor que la unidad: $\rho(K) < 1$ (el radio espectral de K es el máximo de los valores propios de K en valor absoluto). Por tanto, si se reemplaza K en la ecuación (70), se obtiene que

$$H_+A^{-1}\epsilon^{n+1} = H_-A^{-1}\epsilon^n. \quad (71)$$

Recordemos que

$$H_+ = \begin{bmatrix} \Phi_a - \Delta t \theta \mathcal{L} \Phi_a \\ \mathcal{B} \Phi_b \end{bmatrix}, \quad H_- = \begin{bmatrix} \Phi_a + (1 - \theta) \Delta t \mathcal{L} \Phi_a \\ 0_b \end{bmatrix}.$$

Si $\mathcal{B} = I$ donde $I \in \mathbb{R}^{N \times N}$ representa la matriz identidad, se tiene que:

$$H_+ = \begin{bmatrix} \Phi_a \\ \Phi_b \end{bmatrix} - \Delta t \theta \begin{bmatrix} \mathcal{L} \Phi_a \\ 0_b \end{bmatrix}, \quad H_- = \begin{bmatrix} \Phi_a \\ 0_b \end{bmatrix} + (1 - \theta) \Delta t \begin{bmatrix} \mathcal{L} \Phi_a \\ 0_b \end{bmatrix}.$$

Por lo tanto, si se sustituye H_+ , H_- en (71) se tiene el siguiente sistema

$$[I - \theta \Delta t M] \epsilon^{n+1} = [I + (1 - \theta) \Delta t M] \epsilon^n,$$

donde $M = [\mathcal{L}\Phi_a \ 0]^T A^{-1}$. Así, la condición estabilidad se puede expresar como

$$\rho([I - \theta \Delta t M]^{-1} [I + (1 - \theta) \Delta t M]) < 1, \quad (72)$$

es decir,

$$\left| \frac{1 + (1 - \theta) \Delta t \lambda_M}{1 - \theta \Delta t \lambda_M} \right| \leq 1 \quad (73)$$

donde $\lambda_M = \max_{1 \leq i \leq N} \lambda_M^i$ es el máximo autovector asociado a la matriz M .

De la ecuación (73) se tiene que, para $\theta \geq 0.5$ y $\lambda_M \leq 0$, el método implícito es incondicionalmente estable. Por otra parte, cuando $\theta = 0$ se obtiene que el método explícito es condicionalmente estable si

$$|1 + \Delta t \lambda_M| \leq 1. \quad (74)$$

De la ecuación (74) se puede ver que cuando $\lambda_M > 0$, el método explícito es estable siempre que

$$\Delta t \leq 2/\lambda_M.$$

Tanto para el método implícito como para el explícito se necesita que la matriz del sistema sea no singular y esté bien condicionada.

3.3. ECUACIÓN LINEAL DEL CALOR

El objetivo de esta sección es resolver un problema parabólico en dos dimensiones mediante el método de colocación asimétrico. Consideramos la ecuación del calor en un dominio rectangular con condición de frontera Dirichlet homogénea. Se describirá con detalle los esquemas de discretización temporal; implícito y explícito vistos en la Sección 3.1 para el problema mencionado. Además se analizará experimentalmente la tasa de convergencia empleando los esquemas antes mencionados.

Consideremos el siguiente problema lineal evolutivo en un dominio rectangular $\Omega = (0, 1) \times (0, 1)$,

$$\frac{\partial u}{\partial t} - \frac{\partial^2 u}{\partial x^2} - \frac{\partial^2 u}{\partial y^2} = 0, \quad (x, y) \in (0, 1) \times (0, 1). \quad (75)$$

La condición inicial y las condiciones de frontera Dirichlet dependientes del tiempo se consideran utilizando la solución exacta

$$u(x, y, t) = e^{-2\pi^2 t} \sin(\pi x) \sin(\pi y). \quad (76)$$

Inicialmente se describe de forma algorítmica la implementación del sistema iterativo (66) usando el esquema implícito para la ecuación del calor (75). Además se espera determinar experimentalmente la tasa de convergencia empleando el método de colocación asimétrico usando la función radial multicuadrática.

Esquema Implícito

Definamos inicialmente el conjunto de centros X distribuidos uniformemente en 36 nodos en la frontera y 64 nodos en el interior. Con base en los $N = 100$ nodos de colocación se construye la matriz $A = [\Phi_a \ \Phi_b]$ y las matrices $\{\mathcal{L}\Phi_a \ \Phi_b\}$. Para el problema de la ecuación del calor $\mathcal{L}$ denota el operador Laplaciano.

$$\begin{bmatrix} \Phi_a - \frac{\Delta t}{2} \mathcal{L}\Phi_a \\ \Phi_b \end{bmatrix} \lambda^{n+1} = \begin{bmatrix} \Phi_a + \frac{\Delta t}{2} \mathcal{L}\Phi_a \\ 0_b \end{bmatrix} \lambda^n. \quad (77)$$

Por tanto, de (77) se obtiene el sistema

$$H_+ \lambda^{n+1} = H_- \lambda^n, \quad (78)$$

donde

$$H_+ = \begin{bmatrix} \Phi_a - \frac{\Delta t}{2} \mathcal{L} \Phi_a \\ \mathcal{B} \Phi_b \end{bmatrix}, \quad H_- = \begin{bmatrix} \Phi_a + \frac{\Delta t}{2} \mathcal{L} \Phi_a \\ 0_b \end{bmatrix}.$$

Finalmente, del análisis de estabilidad visto en la Sección 3.2, podemos concluir que

$$u^{n+1} = K u^n, \quad (79)$$

siendo $K = A H_+^{-1} H_- A^{-1}$.

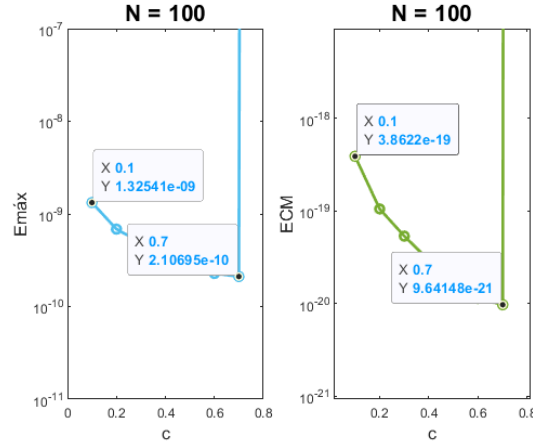
Los pasos necesarios para determinar la solución en el tiempo $t_{\text{máx}}$ del sistema iterativo (79) se muestran en el Algoritmo 1.

Algoritmo 1. Método Implícito

$t = 0$
 Construir Φ_a , Φ_b , $\mathcal{L} \Phi_a$ y Φ_b
 Se define la condición inicial: $u^0 = u(x, y, 0)$
while $t < t_{\text{máx}}$
 $u^{n+1} = K u^n$
 $t = t + \Delta t$
 $u^n = u^{n+1}$
end

Como en el caso del problema lineal estacionario, nos interesa saber si para cualesquiera c y N , el problema de ecuación del calor (75) converge exponencialmente. Analizaremos primero el esquema implícito usando el método de aproximación estacionaria (ver Sección 1.5) con incrementos en c a razón de 0.1, tomando a $N = 100$, $\Delta t = 0.01$, y $t_{\text{máx}} = 1$, ver Figura 13.

Figura 13. Método de aproximación estacionaria: El número de nodos se fija en $N = 100$, $\Delta t = 0.01$, y $t_{\text{máx}} = 1$. Izquierda: Error máximo frente al parámetro de forma c . Derecha: Error cuadrático medio vs parámetro de forma c .



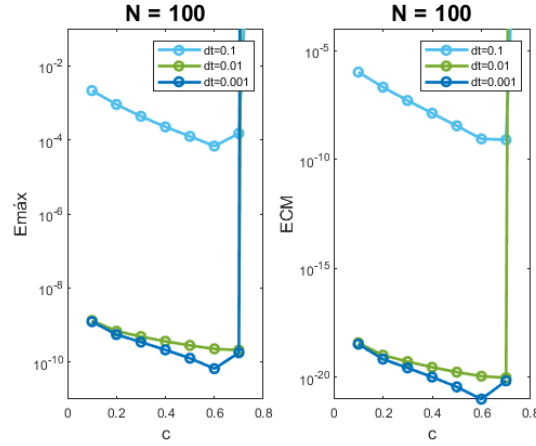
Los resultados (ver Tabla 5) muestran que el método implícito se comporta bien comparado con el caso estacionario, en realidad, el método implícito reduce el error para $c \leq 0.7$. Sin embargo, cuando $c > 0.7$ la matriz de interpolación del sistema se vuelve mal condicionada y los errores crecen exponencialmente.

Tabla 5. Variación del parámetro c para el esquema implícito con $N = 100$, $\Delta t = 0.01$, $t_{\text{máx}} = 1$.

Índice	c	ECM	$E_{\text{máx}}$
1.	0.1	3.8622×10^{-19}	1.3254×10^{-9}
2.	0.2	1.0440×10^{-19}	6.8531×10^{-10}
3.	0.3	5.3140×10^{-20}	4.8953×10^{-10}
4.	0.4	2.9134×10^{-20}	3.6326×10^{-10}
5.	0.5	1.7382×10^{-20}	2.8137×10^{-10}
6.	0.6	1.1303×10^{-20}	2.2757×10^{-10}
7.	0.7	9.6415×10^{-21}	2.1069×10^{-10}

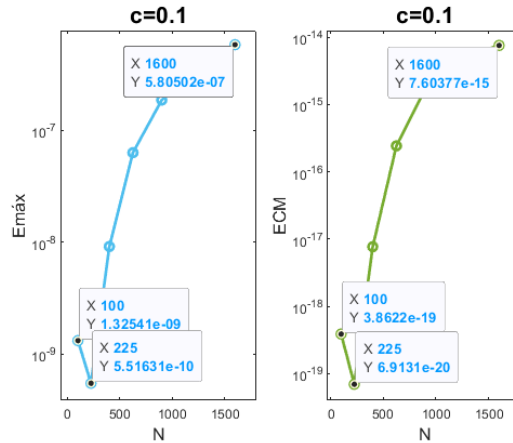
Por otra parte, se observó que el decrecimiento exponencial en el error con el esquema implícito está relacionado con el tamaño del paso en el tiempo Δt . En efecto, al incrementar Δt se incrementa el error de la aproximación, ver Figura 14.

Figura 14. Método de aproximación estacionaria: El número de nodos se fija en $N = 100$, $t_{\text{máx}} = 1$ y distintos pasos temporales Δt . Izquierda: Error máximo frente al parámetro de forma c . Derecha: Error cuadrático medio vs parámetro de forma c .



Ahora hacemos un análisis del esquema implícito usando el método no estacionario para los siguientes valores $\Delta t = 0.01$, $t_{\text{máx}} = 1$, $c = 0.1$ y nodos distribuidos uniformemente definidos en el rango $N = \{100, 225, 400, 625, 900, 1225, 1600\}$, ver Figura 15

Figura 15. Método de aproximación no estacionaria: El parámetro de forma se fija en $c = 0.1$, $\Delta t = 0.01$ y $t_{\text{máx}} = 1$. Izquierda: Error máximo frente al número de centros N . Derecha: Error cuadrático medio vs número de centros N .



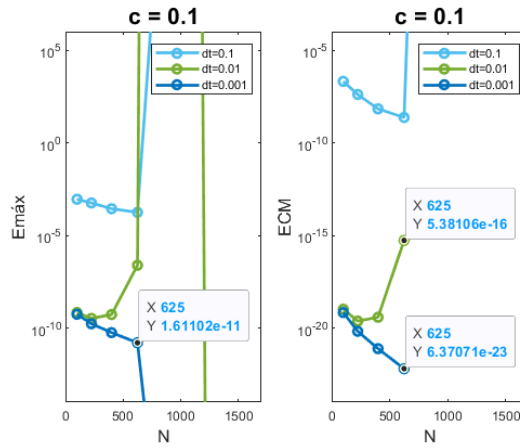
Los resultados (ver Tabla 6) muestran que los errores ECM y $E_{\text{máx}}$ se incrementan conforme crece N . Sin embargo, pudimos observar al igual que el método estacionario una reducción del tamaño del paso temporal Δt genera mejor resultado en la precisión.

Tabla 6. Variación del número de nodos N para el esquema implícito con $c = 0.1$, $\Delta t = 0.01$, $t_{\text{máx}} = 1$.

Índice	N	ECM	$E_{\text{máx}}$
1.	100	3.8622×10^{-19}	1.3254×10^{-9}
2.	225	6.9131×10^{-20}	5.5163×10^{-10}
3.	400	7.6880×10^{-18}	9.1711×10^{-9}
4.	625	2.4357×10^{-16}	6.3098×10^{-8}
5.	900	1.7310×10^{-15}	1.8599×10^{-7}
6.	1225	4.9812×10^{-15}	3.4172×10^{-7}
7.	1600	7.6038×10^{-15}	5.8050×10^{-7}

La Figura 16 expone que la variación del tamaño temporal Δt produce un efecto similar al método estacionario visto en la Figura 14. Por otro lado, si se observa la figura de la derecha de (16), para $\Delta t = 0.001$ y $\Delta t = 0.001$ el error ECM se calcula hasta los $N = 625$. Esto se debe al mal condicionamiento del sistema cuando N se incrementa.

Figura 16. Método de aproximación no estacionaria: El parámetro de forma se fija en $c = 0.1$, $t_{\text{máx}} = 1$ y distintos pasos temporales Δt . Izquierda: Error máximo frente al número de centros N . Derecha: Error cuadrático medio vs número de centros N .



Esquema Explícito Definamos inicialmente el conjunto de centros X distribuidos uniformemente en 36 nodos en la frontera y 64 en el interior, con base en los $N = 100$ nodos de colocación se construye la matriz $A = [\Phi_a \ \Phi_b]$ y las matrices $\{\mathcal{L}\Phi_a \ \Phi_b\}$, para

el problema de la ecuación de calor $\mathcal{L}$ denota el operador Laplaciano.

$$\begin{bmatrix} \Phi_a \\ \Phi_b \end{bmatrix} \lambda^{n+1} = \begin{bmatrix} \Phi_a - \Delta t \mathcal{L} \Phi_a \\ 0_b \end{bmatrix} \lambda^n. \quad (80)$$

Por tanto, de (77) se obtiene un sistema

$$H_+ \lambda^{n+1} = H_- \lambda^n \quad (81)$$

donde

$$H_+ = \begin{bmatrix} \Phi_a \\ \Phi_b \end{bmatrix}, \quad H_- = \begin{bmatrix} \Phi_a - \Delta t \mathcal{L} \Phi_a \\ 0_b \end{bmatrix}.$$

Finalmente, del capítulo de análisis de estabilidad 3.2, podemos concluir que

$$u^{n+1} = K u^n \quad (82)$$

siendo $K = A H_+^{-1} H_- A^{-1}$.

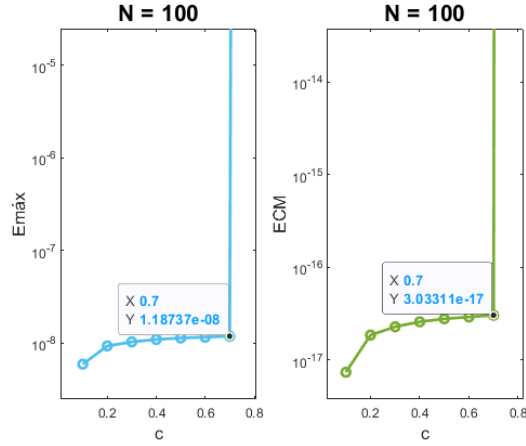
Los pasos necesarios para determinar la solución al tiempo $t_{\text{máx}}$ del sistema iterativo (82) se muestran en el Algoritmo 2.

Algoritmo 2. Método Explícito

$t = 0$
 Construir Φ_a , Φ_b , $\mathcal{L}\Phi_a$ y Φ_b
 Se define la condición inicial: $u^0 = u(x, y, 0)$
while $t < t_{\text{máx}}$
 $u^{n+1} = K u^n$
 $t = t + \Delta t$
 $u^n = u^{n+1}$
end

A continuación se analiza el esquema explícito del mismo modo que el esquema implícito, se resolvió la aproximación numérica usando el método estacionario y no estacionario. Primero, defina los siguientes valores $N = 100$, $\Delta t = 0.01$, $t_{\text{máx}} = 1$, ver Figura 17.

Figura 17. Método de aproximación estacionaria: El número de nodos se fija en $N = 100$, $t_{\text{máx}} = 1$ y $\Delta t = 0.01$. Izquierda: Error máximo frente al parámetro de forma c . Derecha: Error cuadrático medio vs parámetro de forma c .



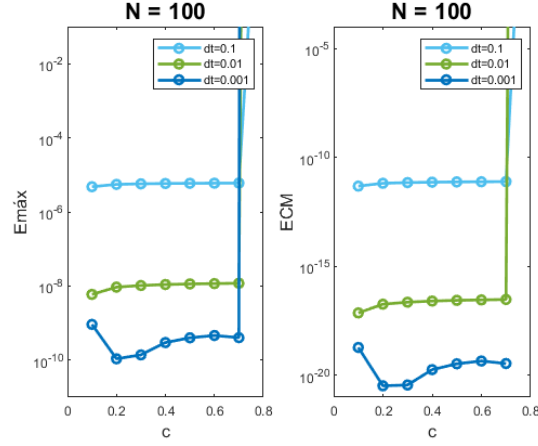
El esquema explícito a diferencia del implícito tiene una tasa de crecimiento lineal sobre el error puntual para $c \leq 0.7$. Sin embargo, el condicionamiento de la matriz se ve afectado para valores de $c > 0.7$. Los resultados presentados en las Tablas 5, 7; verifican la eficiencia del esquema implícito sobre el explícito.

Tabla 7. Variación del parámetro c para el esquema explícito con $N = 100$, $\Delta t = 0.01$, $t_{\text{máx}} = 1$.

Índice	c	ECM	$E_{\text{máx}}$
1.	0.1	7.3502×10^{-18}	5.9351×10^{-9}
2.	0.2	1.8486×10^{-17}	9.3342×10^{-9}
3.	0.3	2.2778×10^{-17}	1.0330×10^{-8}
4.	0.4	2.5749×10^{-17}	1.0964×10^{-8}
5.	0.5	2.7763×10^{-17}	1.1373×10^{-8}
6.	0.6	2.9094×10^{-17}	1.1634×10^{-8}
7.	0.7	3.0331×10^{-17}	1.1874×10^{-8}
8.	0.8	1.0935×10^{149}	6.3945×10^{74}

Si repetimos el procedimiento de variar el tamaño del paso del tiempo Δt en el esquema explícito se observan las mismas consecuencias sobre la precisión cuando se trabaja en el esquema implícito, vea la Figura 18.

Figura 18. Método de aproximación estacionaria: El número de nodos se fija en $N = 100$, $t_{\text{máx}} = 1$ y distintos pasos temporales Δt . Izquierda: Error máximo frente al parámetro de forma c . Derecha: Error cuadrático medio vs parámetro de forma c .



Para el caso no estacionario, el esquema explícito mejora la precisión conforme disminuye el paso temporal Δt , ver Figura 19. Sin embargo, conforme crece N se pierde la convergencia exponencial.

Cabe recordar que en el esquema de colocación asimétrico Cheng ²² muestra numéricamente que “incrementando el número de nodos N sin el correspondiente cambio en el parámetro c , es equivalente en mantener fijo N e incrementar el valor de c ”, según lo anterior, si modificamos el parámetro de forma $c = 0.5$, en la figura izquierda de (20) se observa que la convergencia exponencial se conserva, mientras que para la figura derecha de (20) se tiene un desbordamiento numérico, como consecuencia del mal condicionamiento del sistema algebraico.

Figura 19. Método de aproximación no estacionaria: El parámetro de forma se fija en $c = 0.1$, $t_{\text{máx}} = 1$ y distintos pasos temporales Δt . Izquierda: Error máximo frente al número de nodos N . Derecha: Error cuadrático medio vs número de nodos N .

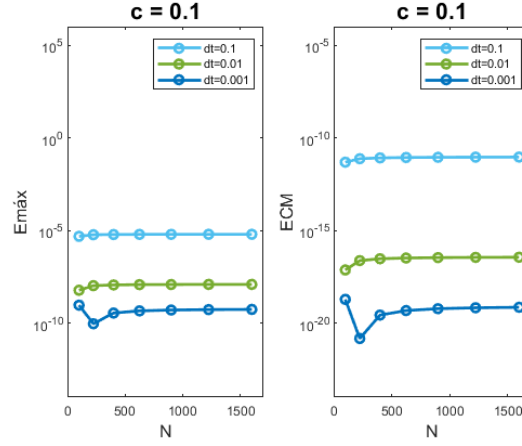
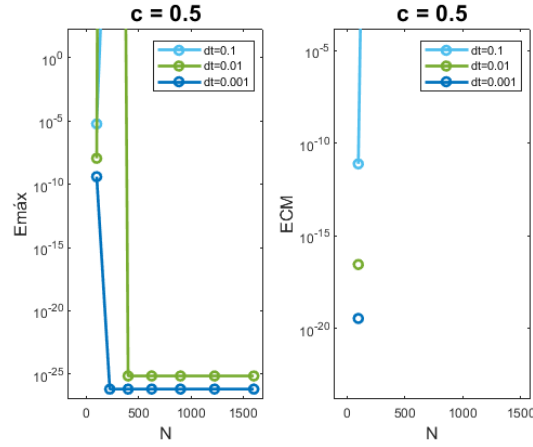


Figura 20. Método de aproximación no estacionaria: El parámetro de forma se fija en $c = 0.5$, $t_{\text{máx}} = 1$ y distintos pasos temporales Δt . Izquierda: Error máximo frente al número de nodos N . Derecha: Error cuadrático medio vs número de nodos N .



3.4. COMPARACIÓN ENTRE EL MÉTODO ELEMENTOS FINITOS Y MÉTODO DE COLOCACIÓN ASIMÉTRICO

En una forma similar a lo realizado en la sección 2.3 se pretende verificar la eficiencia del método de colocación asimétrico sobre los métodos clásicos de aproximación numérica, especialmente el método de elementos finitos. Para este propósito, se emplean el método elementos finitos y método de colocación asimétrico sobre el problema lineal evolutivo (75) con el fin de comparar las respectivas soluciones numéricas.

La solución de la ecuación de calor (75) mediante elementos finitos se puede visualizar en la Figura 22. La implementación numérica es realizada mediante la función *pdetool* de Matlab la cual nos da por defecto un mallado de $N = 1541$.

Figura 21. Dominio de la superficie discretizado en 1541 nodos.

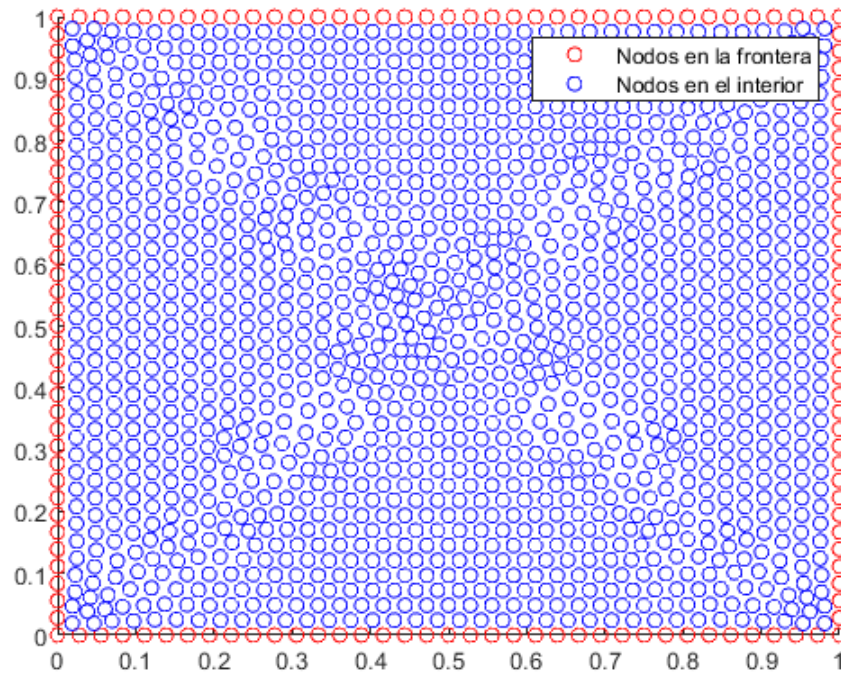
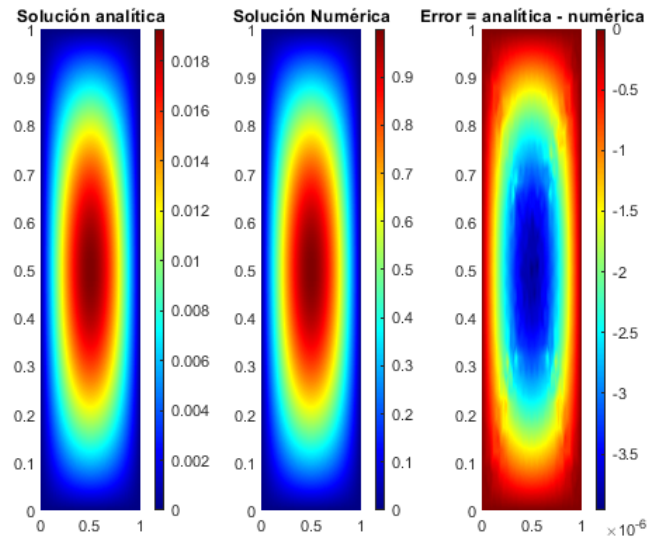


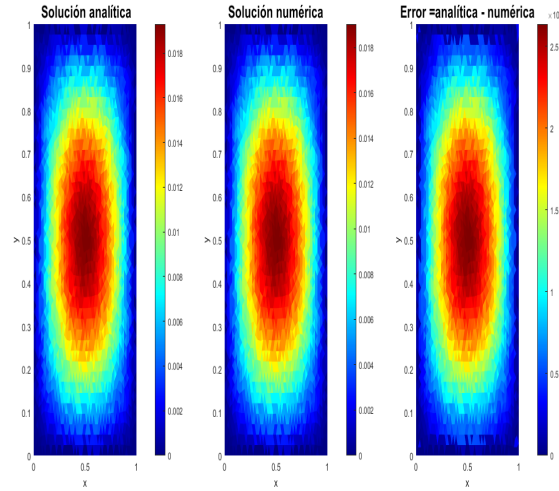
Figura 22. Solución numérica del problema a través del método de elementos finitos. Izquierda: Solución analítica ecuación del calor (75). Centro: Solución numérica de ecuación de calor (75). Derecha: Error puntual entre la solución analítica y solución numérica del problema (75).



La implementación numérica se efectuó con parámetros $t_{\text{máx}} = 0.2$, $\Delta t = 0.01$. Los errores alcanzados por la simulación son $ECM = 4.4495 \times 10^{-8}$ y $E_{\text{máx}} = 5.4 \times 10^{-3}$. La implementación numérica tiene un tiempo de cómputo de $cputime = 36.739705$ segundos.

Ahora nos interesa solucionar la ecuación del calor (75) mediante el método colocación asimétrico tomando como centros de colocación los 1541 nodos de la malla creada para el método de elementos finitos. La implementación se realiza bajo las mismas condiciones que en elementos finitos adicionando el parámetro de forma $c = 0.1$, ver (23). Los errores producidos por la implementación del método colocación asimétrico son $ECM = 1.6629 \times 10^{-8}$ y $E_{\text{máx}} = 2.6348 \times 10^{-4}$. La implementación numérica tiene un tiempo de cómputo de $cputime = 23.427289$ segundos.

Figura 23. Solución numérica método de colocación asimétrico. Izquierda: Solución analítica ecuación de calor (75). Centro: Solución numérica de ecuación de calor (75). Derecha: Error puntual entre la solución analítica y solución numérica del problema (75).



Note que aunque el método colocación asimétrico muestra mejor resultado que elementos finitos, la diferencia no es sustancialmente significativa. Recordemos que en el método colocación asimétrico un refinamiento de la malla o una buena elección del parámetro de forma nos proporcionaría mejor precisión, sin embargo, ambos nos generan mal condicionamiento en el sistema y un alto costo computacional. En la sección anterior, pudimos visibilizar la relación existente entre el tamaño del tiempo y el parámetro de forma en el esquema implícito, en realidad se observó que una variación del tamaño del tiempo produce un efecto similar al variar el parámetro de forma.

En la Tabla 8 mostramos el efecto de variar el tamaño del tiempo Δt , bajo las mismas condiciones iniciales, es decir, $N = 1541$, $t_{\text{máx}} = 0.2$, y $c = 0.1$. Los resultados de la Tabla 8 son satisfactorios; efectivamente verifica la convergencia exponencial del método de colocación asimétrico y la superioridad de éste respecto al método clásico de elementos finitos. Adicionalmente, en la Tabla 8 mostramos el tiempo de cómputo gastado por cada método para diferentes tamaños del tiempo.

Tabla 8. Comparación entre la variación del tamaño temporal en el método elementos finitos, y método de colocación asimétrico (esquema implícito), la disminución de los errores ECM , $E_{\text{máx}}$ y el tiempo de cómputo de la implementación para ambos métodos numéricos.

Método	Δt	ECM	$E_{\text{máx}}$	cputime
M. Colocación asimétrico	0.005	1.5237×10^{-9}	7.8123×10^{-5}	37.509030 s
M. Elementos finitos	0.005	4.5438×10^{-8}	1.11×10^{-2}	66.154538 s
M. Colocación asimétrico	0.0025	2.7539×10^{-10}	3.1755×10^{-5}	46.194469 s
M. Elementos finitos	0.0025	4.6135×10^{-8}	2.28×10^{-2}	118.143013 s
M. Colocación asimétrico	0.002	1.9402×10^{-10}	2.6225×10^{-5}	70.690055 s
M. Elementos finitos	0.002	4.6263×10^{-8}	2.86×10^{-2}	153.583085 s

BIBLIOGRAFÍA

- Chang, K. F. "Strictly positive definite functions". En: *J. Approx. Theory* 87 (1996), págs. 148-158 (vid. pág. 24).
- Cheng H.D., Golberg M. A. Kansa E. J. y G. Zammito. "Exponential convergence and h-c multiquadric collocation method for partial differential equations," en: *Numerical Methods Partial Differential Equations*, 19(5) (2003), págs. 571-594 (vid. págs. 52, 78).
- Cheng, H.D. "Multiquadric and its shape parameter-A numerical investigation of error estimate, condition number, and round-off error by arbitrary precision computation," en: *Engineering Analysis with Boundary Elements* 36 (2012), págs. 220-239 (vid. págs. 48, 50).
- Chinchapatnam P. P, Djidjeli K. y P.B. Nair. "Domain decomposition for time-dependent problems using radial based meshless methods." En: *Numerical Methods for Partial Differential Equations*, 23(1) (2007), págs. 38-59 (vid. pág. 68).
- Ciarlet, G. *Introduction to Numerical Linear Algebra and Optimisation* (vid. pág. 90).
- De Marchi. S, Schaback. R y Wendland. H. "Near-optimal data-independent point locations for radial basis function interpolation," en: *Advances in Computational Mathematics*, 23 (2005), págs. 317-330 (vid. pág. 64).
- Duchon, J. "Sur l'erreur d'interpolation des fonctions de plusieurs variables par les D^m -splines," en: *Rev. Française Automat. Informat. Rech. Opér. Anal.* 12 (1978), págs. 325-334 (vid. pág. 38).
- Fasshauer, G. E. *Boundary Elements XXVII: Incorporating Electrical Engineering and Electromagnetics, chapter RBF Collocation Methods and Pseudospectral Methods*, (vid. pág. 13).

- Fasshauer, G. E. *Meshfree Approximation Methods with MATLAB*, World Scientific, (vid. págs. 19, 21, 23, 26, 28, 29, 37, 38, 42, 43).
- Foley, T. A. "Near optimal parameter selection for multiquadric interpolation," en: *Journal of Applied Science and Computation*, 1 (1994), págs. 54-69 (vid. pág. 51).
- Fornberg, B. y G. Wright. "Stable computation of multiquadric interpolants for all values of the shape parameter," en: *Computers and Mathematics with applications*, 48 (2004), págs. 853-867 (vid. pág. 64).
- Franke, R. "Scattered data interpolation: test of some methods," en: *In Mathematics of Computation*, 38 (1982), págs. 181-200 (vid. págs. 48, 51).
- Golberg, M.A. y C.S. Chen. *Boundary Integral Methods - Numerical and Mathematical Aspects, chapter The method of fundamental solutions for potential, Helmholtz and diffusion problems, pages 103–176*, (vid. pág. 13).
- Hardy, R. L. "Multiquadric equations of topography and other irregular surfaces," en: *Journal of Geophysical Research*, 76(8) (1971), págs. 1905-1915 (vid. pág. 51).
- Higham, N. J. *Accuracy and Stability of Numerical Algorithms*, (vid. pág. 90).
- Huang C.S, Leeb C.F. y A.H.D. Cheng. "Error estimate, optimal shape factor, and high precision computation of multiquadric collocation method," en: *Engineering Analysis with Boundary Elements*, 31 (2007), págs. 614-623 (vid. pág. 64).
- Kansa, E. "Multiquadrics—a scattered data approximation scheme with applications to computational fluid-dynamics. II. solutions to parabolic, hyperbolic and elliptic partial differential equations," en: *Computers and Mathematics with Applications*, 19(8-9) (1990), 147–161 (vid. págs. 10, 11, 13, 53).
- Kansa, E. J. "Multiquadrics—a scattered data approximation scheme with applications to computational fluid-dynamics. I. surface approximations and partial deriva-

- tive estimates,” en: *Computers and Mathematics with Applications*, 19(8-9) (1990), págs. 127-145 (vid. págs. 10, 11, 13, 53).
- Kansa, E. J. y R. E. Carlson. “Improved accuracy of multiquadric interpolation using variable shape parameter,” en: *Computers and Mathematics with Applications*, 24 (1992), págs. 99-120 (vid. pág. 52).
- Li, J. e Y. C. Hon. “Domain decomposition for radial basis meshless methods,” en: *Numerical Methods for Partial Differential Equations*, 20(3) (2004), págs. 450-462 (vid. pág. 64).
- Ling, L. y Kansa. E.J. “A least squares preconditioner for radial basis functions collocation methods,” en: *Advances in Computational Mathematics*, 23 (2005), págs. 31-54 (vid. pág. 64).
- Ling, L. y E.J. Kansa. “Preconditioning for Radial Basis Functions with Domain Decomposition Methods,” en: *Mathematical and Computer Modelling, Submitted to Mathematical and Computer Modeling*, 40 (2004), págs. 1413-1427 (vid. pág. 64).
- Madych W. R., Nelson S. A. “Polyharmonic cardinal splines,” en: *Journal of Approximation Theory*, 60 (1990), págs. 141-156 (vid. pág. 38).
- Madych, W. R. “Miscellaneous error bounds for multiquadric and related interpolato,” en: *Computers and Mathematics with Applications*, 24(12) (1992), págs. 121-138 (vid. pág. 49).
- Madych, W. R. y S. A. Nelson. “Bounds on multivariate polynomials and exponential error estimates for multiquadric interpolation,” en: *Journal of Approximation Theory*, 70(1) (1992), págs. 94-114 (vid. pág. 48).
- Micchelli, C.A. “Interpolation of scattered data-distance matrices and conditionally positive definite functions,” en: *Constructive Approximation*, 2(1) (1986), págs. 11-22 (vid. págs. 30, 31).

- Muñoz, J. *Solución numérica de ecuaciones en derivadas parciales mediante funciones radiales (tesis doctoral)*, (vid. págs. 13, 90).
- Schaback, R. "Error estimates and condition numbers for radial basis function interpolation," en: *Advances in Computational Mathematics*, 3 (1995), págs. 251-264 (vid. pág. 50).
- Tolstykh, A. I. y D. A. Shirobokov. "On using radial basis functions in a finite difference mode with applications to elasticity problems," en: *Computational Mechanics*, 33(1) (2003), págs. 68-79 (vid. pág. 12).
- Trahan, C. J. y R. E. Wyatt. "Radial basis function interpolation in the quantum trajectory method: optimization of the multiquadric shape parameter," en: *Journal of Computational Physics*, 185 (2003), págs. 27-49 (vid. pág. 51).
- Trefethen, L. N. y D. Bau. *Numerical Linear Algebra, first edition*, (vid. pág. 64).
- Wendland, H. "Meshless Galerkin methods using radial basis functions," en: *Mathematics of Computatio*, 68(228) (1999), págs. 1521-1531 (vid. pág. 13).
- Wendland, Holguer. *Scattered Data Approximation, Cambridge Monographs on Applied and Computational Mathematics*, (vid. págs. 24, 30, 32-35, 37, 40, 46).
- Wright, G. *Radial Basis Function Interpolation: Numerical and Analytical Developments, PhD thesis*, (vid. pág. 64).
- Wu, Z. y R. Schaback. "Local error estimates for radial basis function interpolation of scattered data," en: *IMA Journal of Numerical Analysis*, 13(1) (1993), págs. 13-21 (vid. págs. 38, 39).
- Zerroukat M, Djidjeli. K y Charafi. "Explicit and implicit meshless methods for linear advection-diffusion-type partial differential equations," en: *International Journal of Numerical Methods in Engineering*, 48(1) (2000), págs. 19-35 (vid. pág. 68).

4. Número de condición de una matriz

Como es bien sabido, el número de condición de una matriz juega un papel importante en la resolución de sistemas lineales, especialmente en los problemas de perturbación. El propósito de esta sección es conocer algunas de sus propiedades y definir el significado de mal condicionamiento. Consideremos el sistema lineal de ecuaciones

$$Ax = b, \quad (83)$$

donde b es un vector en $\mathbb{R}^d$, A es una matriz real de dimensiones $d \times d$ y x es el vector de incógnitas en $\mathbb{R}^d$. Para determinar la solución del sistema lineal (83), es necesario recordar que este sistema tiene única solución si, y solo si, la matriz A es no singular, es decir, si $\det(A) \neq 0$. En este caso, $x = A^{-1}b$, donde A^{-1} es la matriz inversa de A .

El aspecto numérico que abordaremos es el problema de la estabilidad de los sistemas lineales de ecuaciones. Nos interesa saber qué ocurre con la solución al realizar una pequeña perturbación en los datos. Es importante considerar la cuantificación del error absoluto y el error relativo de tales perturbaciones.

Introduciendo una perturbación δx en x y una perturbación δb en b , el sistema (83) queda expresado como

$$A(x + \delta x) = b + \delta b, \quad (84)$$

es decir,

$$\begin{aligned} \delta x &= A^{-1}b + A^{-1}\delta b - A^{-1}Ax \\ &= A^{-1}b + A^{-1}\delta b - x \\ &= x + A^{-1}\delta b - x \\ &= A^{-1}\delta b. \end{aligned} \quad (85)$$

Usando las propiedades de las normas entre matrices, de (85) se obtiene

$$\|\delta x\| \leq \|A^{-1}\| \cdot \|\delta b\|. \quad (86)$$

Esta desigualdad acota el cambio absoluto de la solución en función del cambio absoluto del vector conocido b . Un inconveniente de utilizar el error absoluto es que no considera el orden en la magnitud del valor que se estima. Es por ello que se introduce el error relativo, el cual lo podemos expresar como un porcentaje del error cometido independientemente de las magnitudes que estén midiendo.

Ahora, lo que nos interesa es comparar las cantidades del error relativo entre las perturbaciones de los vectores $x, b \in \mathbb{R}^d$, donde el error relativo se expresa

$$\frac{\|x - \delta x\|}{\|x\|}, \frac{\|b - \delta b\|}{\|b\|}. \quad (87)$$

Tomando (84) y computando su norma, se obtiene

$$\|x - \delta x\| \leq \|A^{-1}\| \cdot \|b - \delta b\|. \quad (88)$$

Por otra parte, como $b = Ax$, entonces $\|b\| = \|Ax\|$, por tanto, $\frac{1}{\|x\|} \leq \frac{\|A\|}{\|b\|}$, lo que al combinarlo con la ecuación (88) da como resultado

$$\frac{\|x - \delta x\|}{\|x\|} \leq \|A\| \cdot \|A^{-1}\| \frac{\|b - \delta b\|}{\|b\|}. \quad (89)$$

El número $\text{cond}(A) = \|A\| \cdot \|A^{-1}\|$ se llama número de condición de A . A menor número de condición decimos que el sistema está bien condicionado, a mayor número de condición el sistema está mal condicionado y pequeños cambios relativos nos inducen a un degradamiento en la solución. Algunas propiedades del número de condición de una matriz se enumeran a continuación

Proposición 4.1. Sea $\|\cdot\|$ la norma matricial y A una matriz invertible. Entonces,

1. $\text{cond}(A) \geq 1$.

$$2. \text{cond}(A) = \text{cond}(A^{-1}).$$

$$3. \text{cond}(\lambda A) = \text{cond}(A) \text{ para todo } \lambda \in \mathbb{R} - \{0\}.$$

Proposición 4.2. Sea A una matriz invertible. Entonces

$$\text{cond}_2(A) = \sqrt{\frac{\lambda_n(A * A)}{\lambda_1(A * A)}}$$

donde $\text{cond}_2(A) = \|A\|_2 \|A^{-1}\|_2$, $\|A\|_2$ denota la norma espectral correspondiente a la norma vectorial euclidiana $\|\cdot\|$. Además $\lambda_n(A * A)$ y $\lambda_1(A * A)$ son los autovalores mayor y menor de la matriz $A * A$, respectivamente.

Lema 4.3. Sea A una matriz hermitiana no singular, con autovalores $\lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_n$. Entonces

$$\text{cond}_2(A) = \frac{\lambda_n}{\lambda_1}.$$

Observación 4.4. Sea $\|\cdot\|$ la norma matricial subordinada y A una matriz hermitiana. Entonces

$$\text{cond}(A) = \|A\| \cdot \|A^{-1}\| \geq \text{cond}_2(A).$$

Por tanto, para matrices hermitianas, cond_2 es el menor de todos.

Los anteriores resultados son obtenidos de³³

³³ G. Ciarlet. *Introduction to Numerical Linear Algebra and Optimisation*; N. J. Higham. *Accuracy and Stability of Numerical Algorithms*, J. Muñoz. *Solución numérica de ecuaciones en derivadas parciales mediante funciones radiales (tesis doctoral)*,