

**TÉCNICA DE REGULARIZACIÓN PARA LA IDENTIFICACIÓN
SIMULTÁNEA DEL TÉRMINO FUENTE Y LA TEMPERATURA EN LA
ECUACIÓN DEL CALOR**

JENIFER PERALTA NAVARRO

Universidad Del Valle
Facultad de Ciencias Naturales y Exáctas
Departamento de Matemáticas
Santiago de Cali
2012

**TÉCNICA DE REGULARIZACIÓN PARA LA IDENTIFICACIÓN
SIMULTÁNEA DEL TÉRMINO FUENTE Y LA TEMPERATURA EN LA
ECUACIÓN DEL CALOR**

JENIFER PERALTA NAVARRO

Trabajo de investigación presentado para optar al título de Magister en
ciencias-matemáticas

Directora
Doris Hinestroza G., Ph.D.

Universidad Del Valle
Facultad de Ciencias Naturales y Exáctas
Departamento de Matemáticas
Santiago de Cali
2012

Resumen

En este trabajo consideraremos el problema de identificar el término fuente en la ecuación del calor. Esto es, identificar la función $F(x, t, u)$ en

$$u_t(x, t) = u_{xx}(x, t) + F(x, t, u), \quad 0 \leq x \leq 1, \quad 0 \leq t \leq 1, \quad (2.7)$$

bajo las condiciones

$$\begin{aligned} u(x, 0) &= f(x) \quad 0 \leq x \leq 1, \\ u(1, t) &= g(t) \quad 0 \leq t \leq 1, \\ u(0, t) &= h(t) \quad 0 \leq t \leq 1. \end{aligned} \quad (2.8)$$

Con la información adicional que proviene de las medidas

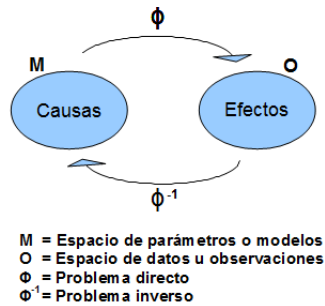
$$u^\delta(x, t) = u(x, t) + \text{ruido}, \quad t \in (t_1, t_2),$$

donde $\|u - u^\delta\| \leq \delta$. Este problema de identificar F es en general no lineal, además como no conocemos la solución $u(x, t)$, en el proceso de aproximar F se identificará también $u(x, t)$.

Este tipo de problemas son mal formulados debido a que involucran derivadas y de haber pequeños errores en los datos generarán grandes errores en los resultados. Para tratar este problema se mostrará un algoritmo de aproximación a F que consiste en minimizar un funcional especial (*teorema principal*) que depende de dos parámetros de regularización α y ϵ , esta técnica de regularización [7] es conocida como regularización generalizada de Tikhonov y para esto debemos introducir las ecuaciones de sensibilidad y la ecuación del problema adjunto.

Introducción

Distinguir entre un problema directo y un problema inverso de un fenómeno dado es en general una controversia, la definición de cada una de ellas está ligado a nuestra cultura, es decir, es lo que interpretamos como una causa-efecto. El investigador ruso **Oleg Mikailivitch Alifanov** en el área de problemas inversos, afirma que: " *una solución de un problema inverso consiste en determinar causas basándose en observaciones de sus efectos*", desde el punto de vista práctico generalmente se llama problema directo aquel en que su estudio fué precedido históricamente. La ambigüedad que genera (directo e inverso) puede ser ejemplificada de la siguiente manera:



El astrofísico georgiano **Viktor Amazaspovich Ambartsumian** en (1996) dio una definición muy amplia de problema inverso: *Un problema inverso consiste en la determinación de causas desconocidas a partir de efectos deseados u observados*, nótese que en general estas observaciones son imprecisas (datos con ruido ó errores experimentales) e incompletas, caso contrario con los problemas directos que requieren un acontecimiento completo y preciso de las causas para determinar los efectos.

A principios del siglo pasado (1902) el matemático francés Jacques Hadamard, introdujo el concepto de problema bien puesto. Según Hadamard, un problema es llamado bien puesto, si la ecuación $Kx = y$ con $K : X \rightarrow Y$ (lineal o no lineal) y X, Y espacios normados cumple que:

1. Existencia: $\forall y \in Y$ existe al menos un $x \in X$ tal que $Kx = y$.
2. Unicidad: $\forall y \in Y$ existe a lo más un $x \in X$ tal que $Kx = y$.

-
3. Estabilidad: La solución x depende continuamente de y , es decir, para toda sucesión $x_n \subset X$ con $Kx_n \rightarrow Kx$ cuando $n \rightarrow \infty$, se sigue que $x_n \rightarrow x$ cuando $n \rightarrow \infty$.

Los problemas que no cumplan al menos una de estas propiedades son llamados problemas mal puestos (ill posed).

Los problemas directos son llamados problemas mal condicionados si no se cumple la condición (3). En general ninguna de las condiciones de Hadamard se cumplen en un problema inverso.

El modelamiento matemático de muchos fenómenos físicos involucran usualmente ecuaciones diferenciales (ordinarias o parciales). Estos fenómenos dan lugar a la aparición de problemas inversos cuya naturaleza es no lineal. Para determinar un modelo matemático completamente se requiere conocer los coeficientes o los términos fuente que determinan el sistema y que están relacionados con sus propiedades físicas en conjunto con las condiciones iniciales y/o datos de frontera. Dentro de la rama de los problemas inversos, está la identificación de parámetros de un sistema físico. Cuando uno o varios coeficientes o el término fuente del sistema no son conocidos, es decir no conocemos el modelo, el problema se convierte en un problema inverso llamado **problema de identificación de coeficientes**. Estos problemas están relacionados con la determinación de propiedades físicas mediante observaciones hechas sobre la frontera de una determinada región acotada.

El problema inverso de identificar el término fuente en la ecuación del calor (2.7)-(2.8) es de vital importancia en muchas ramas de las ciencias de la ingeniería. Por ejemplo, una estimación precisa de la fuente contaminante es fundamental para salvaguardar el medio ambiente en ciudades con poblaciones muy altas. En el problema de la conducción de calor y la difusión, siendo u la temperatura, la función $F(u)$ desconocida se interpreta como una fuente de calor y del material considerado. Por otro lado en aplicaciones relacionadas con productos químicos o bioquímicos la función $F(u)$ se puede interpretar como un término de reacción.

El problema (2.7)-(2.8) ha sido estudiado recientemente en diversos trabajos [2,3,5,8]. Para los problemas inversos se han propuesto varios procedimientos de regularización por ejemplo: Cannon y DuChateau [2] considera un ajuste de mínimos cuadrados para analizar el problema inverso del calor, la idea es determinar los parámetros óptimos a fin de minimizar el error en el funcional a partir de los datos con errores. Sin embargo, este método tiene algunos problemas, puesto que no suele ser evidente que la solución del problema de optimización solucione el problema inverso original y el funcional del error puede estar basado sobre datos que no determinan únicamente la función desconocida. Otros métodos para resolver este problema son los métodos de actualización residuales tales como Newton, homotopía y Proyección de punto fijo (FPP), la principal dificultad con estos métodos es la forma de la no linealidad [3]. Por otro lado Fatullayev [5] propone una aproximación de la función desconocida por piezas poligonales lineales determinadas

consecutivamente de la solución del problema de minimización asociado. En [8] aplican la técnica de regularización de Tikhonov para formular una solución regularizada, que es estable y converge a la solución exacta con una estimación de error de tipo logarítmico.

Es bueno considerar que en estos trabajos no se han considerado errores en los datos. En los problemas no lineales son muy pocos los análisis que se han hecho, para nuestro trabajo aplicaremos una generalización del método de regularización de Tikhonov [7] donde involucraremos dos parámetros, para esto, se desarrollará el algoritmo de regularización que permita recuperar el término fuente F de (2.7) - (2.8), cuando f, g y h se suponen conocidas y u viene de medidas en la práctica, es decir tiene errores en los datos u^δ . Una de las dificultades matemáticas que se presentará en este problema inverso mal formulado (2.7) - (2.8) es encontrar u_t y u_{xx} debido a que la diferenciación es bien conocida como un problema inverso mal puesto, generalmente la solución no depende continuamente de los datos (inestabilidad), esto es, pequeños cambios en los datos produce grandes cambios en las soluciones computadas, para lo cual, generalmente se recurre a métodos de regularización.

Para desarrollar un algoritmo que nos permita identificar numéricamente el término fuente F , debemos minimizar el funcional de regularización generalizado de Tikhonov, durante este proceso se deben introducir unas ecuaciones especiales conocidas como ecuaciones de sensibilidad y la ecuación del problema del adjunto. Una vez obtenida la minimización del funcional determinaremos unas fórmulas numéricas que se implementarán en el computador para mostrar ejemplos de nuestro interés. Para facilitar la lectura, el trabajo se ha dividido de la siguiente manera:

En el capítulo 1, se presentarán los preliminares con resultados necesarios para afrontar este problema. En el capítulo 2, se hará una introducción a la ecuación del calor como problema directo y a las diferentes aplicaciones y soluciones numéricas. En el capítulo 3, se analizará el método generalizado de Tikhonov, mostrando los teoremas de existencia, convergencia y estabilidad. En el capítulo 4, se desarrollará el método generalizado de Tikhonov al problema considerado en este trabajo para determinar las fórmulas numéricas que nos identifican el término fuente F . Por último en el capítulo 5, se presentará las conclusiones y trabajo a futuro.

Índice general

Introducción	2
1. Preliminares	6
1.1. Problemas inversos y regularización	6
1.2. Molificación para la diferenciación numérica	13
1.2.1. Aproximación de derivadas	13
1.2.2. Diferenciación numérica por molificación discreta	15
1.2.3. Ejemplos	15
1.3. Convergencia débil	17
1.4. Métodos iterativos	19
2. Ecuación del calor	23
2.1. Derivación de las ecuaciones reacción-difusión	23
2.2. Aplicaciones	25
2.3. Teoremas de existencia y unicidad de la solución	28
2.4. Solución numérica	30
2.4.1. Ejemplos	31
3. Regularización Tikhonov Generalizada	35
3.1. Existencia y monotonicidad	36
3.2. Estabilidad y convergencia	41
4. Fórmulas numéricas	46
4.1. Formulación variacional y operadores débilmente cerrados	46
4.2. Ecuaciones de sensibilidad	50
4.3. Ecuación del problema adjunto	52
4.4. Gradiente del funcional	55
4.5. Mejoramiento de la estimación de F	56
4.6. Experimentación numérica	57
5. Conclusiones	65
6. Referencias	67

Capítulo 1

Preliminares

1.1. Problemas inversos y regularización

El astrofísico Viktor Amazaspovich Ambartsumian define problema inverso como determinar *causas desconocidas a partir de efectos deseados u observados*, esta definición la tendremos en cuenta para nuestro trabajo, cabe anotar que dichas observaciones son imprecisas, tenemos errores en los datos debido a que son datos experimentales, serán incompletos, cosa que no sucede con los problemas directos. Para aclarar un poco esta definición de causa-efecto, podemos tener como causas en un modelo matemático, las condiciones iniciales y de frontera, la fuente térmica/sumidero y propiedades (material); los efectos son las propiedades calculadas a partir de un modelo directo, como un campo de temperatura, concentración de partículas, corriente eléctrica, etc.

Los problemas inversos se pueden clasificar de diferentes formas, aquí los clasificaremos de acuerdo a la información que se quiere identificar:

- Problema de retrospectiva: Este problema consiste en determinar las condiciones iniciales del modelo.
- Problema inverso de coeficientes: Este problema es de estimación de parámetros. Por ejemplo identificación del coeficiente de difusión ó fuente en la ecuación del calor.
- Problema inverso de frontera: Se debe encontrar las condiciones de frontera de un dominio. Por ejemplo en el problema inverso de conducción de calor, la temperatura en la frontera es desconocida y se determina a partir de mediciones en el interior del objeto.

El matemático Jacques Hadamard introdujo el concepto de problema bien puesto, para aquellos problemas que cumplen con tres condiciones: existencia, unicidad y estabilidad. Los problemas que no cumplan con al menos una de las tres condiciones son llamados mal puestos.

1. Preliminares

Para ilustrar el concepto de un problema inverso mal puesto, consideremos los siguientes ejemplos, ver [15].

Ejemplo 1. Diferenciación

Definimos como el problema directo el evaluar la integral:

$$(T_D\varphi)(x) := \int_0^x \varphi(t)dt \quad x \in [0, 1],$$

para una función real $\varphi \in C([0, 1])$ dada. El problema inverso consiste en resolver la siguiente ecuación:

$$T_D\varphi = g. \quad (1.1)$$

Supongamos que damos una función $g^\delta \in C^1([0, 1])$ tal que

$$\|g^\delta - g\|_\infty \leq \delta \quad \text{con} \quad 0 < \delta < 1.$$

Las funciones

$$g_n^\delta(x) := g(x) + \delta \sin\left(\frac{nx}{\delta}\right) \quad x \in [0, 1],$$

con $n = 2, 3, 4, \dots$ en particular satisfacen

$$\|g_n^\delta - g\|_\infty = \delta.$$

Sin embargo las derivadas

$$(g_n^\delta)'(x) = g'(x) + n \cos\left(\frac{nx}{\delta}\right) \quad x \in [0, 1],$$

satisfacen

$$\|(g_n^\delta)' - g'\|_\infty = n.$$

Entonces para un error en el dato arbitrariamente pequeño δ , el error en el cálculo de la derivada, es muy grande. Así la derivada es un problema inverso mal puesto.

Ejemplo 2.

Consideremos la ecuación $Ax = b$, $b \in \mathbb{R}^m$, $x \in \mathbb{R}^n$, $A \in \mathbb{R}^{m \times n}$. Sea

$$Z = \{x^* : \|Ax^* - b\| = \inf \{\|Ax - b\| : x \in \mathbb{R}^n\}\},$$

Z : Conjunto de las soluciones del problema de mínimos cuadrados, $R(A)$: Rango de A , $N(A)$: Kernel de A . Se puede demostrar fácilmente que

$$\|Ax^* - b\| = \inf \{\|Ax - b\| : x \in \mathbb{R}^n\} \Leftrightarrow A^T Ax^* = A^T b,$$

1. Preliminares

es decir, x^* es solución de mínimos cuadrados de $Ax^* = b$ si y sólo si Ax^* es el elemento más cercano en $R(A)$ a b , o equivalentemente $Ax^* - b \in R(A)^\perp = N(A)$. Así $A^T(Ax^* - b) = 0$ y por lo tanto $A^T Ax^* = A^T b$.

La solución generalizada $\tilde{x}$ se define como:

$$\|\tilde{x}\| = \inf \{ \|z\| : z \in Z \}.$$

Para ilustrar un poco el ejemplo, veamos que el sistema $Ax = b$, no tiene solución para

$$A = \begin{pmatrix} 1 & 1 \\ 1 & 1 \\ 1 & 1 \end{pmatrix} \quad y \quad b = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}.$$

Así las soluciones de mínimos cuadrados están dadas por

$$Z = \{(x, y) : x + y = 2\},$$

y la solución generalizada está dada por

$$\tilde{x} = \begin{pmatrix} 1 \\ 1 \end{pmatrix}.$$

Ejemplo 3.

Consideremos el espacio de Hilbert ℓ_2 dado por el conjunto de sucesiones reales cuadrado sumables,

$$\ell_2 = \left\{ x = \{x_n\} = (x_1, x_2, \dots, x_n, \dots) : \sum_{n=1}^{\infty} x_n^2 < \infty \right\}.$$

Definamos el operador $A : \ell_2 \rightarrow \ell_2$ con

$$y = Ax = \left(x_1, \frac{x_2}{8}, \frac{x_3}{27}, \dots, \frac{x_n}{n^3}, \dots \right),$$

$$y_n = \frac{x_n}{n^3}, \quad \sum_{n=1}^{\infty} y_n^2 = \sum_{n=1}^{\infty} \frac{x_n^2}{n^6} \leq \sum_{n=1}^{\infty} x_n^2 < \infty.$$

Observemos que el operador A es un operador lineal y acotado puesto que para cada $\alpha, \beta \in \mathbb{R}$

$$A(\alpha\{x_n\} + \beta\{z_n\}) = \alpha A(\{x_n\}) + \beta A(\{z_n\}).$$

Además, puesto que cada elemento de la sucesión $\{1/n^6\}$ es acotada por 1 tenemos que

$$\|Ax\|^2 = \sum_{k=1}^{\infty} \frac{|x_k|^2}{n^6} \leq \sum_{k=1}^{\infty} |x_k|^2 = \|x\|^2,$$

1. Preliminares

es decir

$$\|A\| = \sup_{x \neq 0} \left\{ \frac{\|Ax\|}{\|x\|} \right\} \leq 1.$$

Por otro lado el operador A es compacto ya que podemos definir para cada $k \in \mathbb{N}$ un operador lineal $A_k \in B(\ell_2)$ ¹ tal que $\|A - A_k\| \rightarrow 0$ cuando $k \rightarrow \infty$. En efecto,

$$A_k \{x_n\} = \{c_n^k\} \quad \text{donde} \quad c_n^k = \begin{cases} x_n/n^3 & n \leq k, \\ 0 & n > k. \end{cases}$$

Los operadores A_k son operadores lineales y acotados y tienen rango finito, por lo tanto son compactos y además

$$\|(A - A_k)x\|^2 = \sum_{n=k+1}^{\infty} \frac{|x_n|^2}{n^4} \leq \frac{1}{(k+1)^6} \sum_{k=1}^{\infty} |x_k|^2 = \frac{1}{(k+1)^6} \|x\|^2.$$

De esta desigualdad se sigue que

$$\|A - A_k\| \leq \frac{1}{(k+1)^3},$$

y así $\|A - A_k\| \rightarrow 0$ cuando $k \rightarrow \infty$. Así A es compacto por ser límite de operadores compactos.

Ahora, si deseamos resolver el sistema $Ax = y$ para un y dado donde A es compacto, vemos que no es posible resolver el sistema de la forma $x = A^{-1}y$, esto es debido a que A no es invertible. Si suponemos que lo es, entonces formalmente el operador inverso de A está dado por

$$A^{-1}y = (y_1, 8y_2, 27y_3, \dots, n^3y_n, \dots).$$

Puesto que

$$A^{-1}Ax = A^{-1}A\{x_n\} = A^{-1}\left(x_1, \frac{x_2}{8}, \frac{x_3}{27}, \dots, \frac{x_n}{n^3}, \dots\right) = (x_1, x_2, x_3, \dots, x_n, \dots) = x,$$

y similarmente

$$AA^{-1}y = AA^{-1}\{y_n\} = A(y_1, 8y_2, 27y_3, \dots, n^3y_n, \dots) = (y_1, y_2, y_3, \dots, y_n, \dots) = y,$$

pero el operador A^{-1} no es acotado.

Por último observemos que si tomamos la siguiente sucesión

$$y_n = (0, 0, \dots, 0, \frac{1}{n}, 0, \dots),$$

¹ $B(\ell_2)$: Espacio de las funciones lineales acotadas de ℓ_2 en ℓ_2

1. Preliminares

claramente $y_n \rightarrow 0$, pero $\|A^{-1}y_n\| = n \rightarrow \infty$.

En conclusión si tenemos el sistema $Ax = y$, para un y dado, este sistema no tiene solución para todo $y \in \ell_2$, es decir si

$$y = \left(1, \frac{1}{2^2}, \frac{1}{3^2}, \dots, \frac{1}{n^2}, \dots\right),$$

vemos que

$$y = \left(1, \frac{1}{2^2}, \frac{1}{3^2}, \dots, \frac{1}{n^2}, \dots\right) = A(1, 2, 3, \dots, n, \dots), \text{ pero } (1, 2, 3, \dots) \notin \ell_2 \text{ aunque } y \in \ell_2.$$

Así, este problema no cumple con la condición (1) de Hadamard y se convierte en un problema inverso mal puesto.

Una familia de problemas inversos inherentemente más difíciles son los referidos conjuntamente como problemas inversos no lineales. Estos problemas no lineales tienen una relación más compleja entre los datos y el modelo, aunque también pueden ser representados por la ecuación:

$$y = Ax,$$

donde A es un operador no lineal. En este tipo de problemas, lo primero que se debe hacer es comprender la estructura del problema y dar una respuesta teórica a los cuestionamientos de Hadamard (de tal manera que el problema está resuelto desde el punto de vista teórico). Una vez hecho esto se sigue con el estudio de la regularización y de las interpretaciones de la evolución de las soluciones con nuevas medidas (probabilísticas o de otro tipo).

Para solucionar los diferentes problemas inversos mal puestos, recurrimos a una lista de métodos de soluciones:

- Inversión directa.
- Descomposición en valores singulares.
- Mínimos cuadrados.
- Métodos variacionales.
- Molificación.
- Métodos Bayesianos, etc.

En nuestro trabajo estamos interesados en regularizar nuestro problema, la regularización [17] consiste en la aproximación de un problema mal formulado por una familia cercana de problemas bien formulados, que dependen de un parámetro positivo llamado *parámetro de regularización*. La principal propiedad es que, en el caso de datos libres de ruido,

1. Preliminares

las funciones de la familia convergen exactamente a la solución del problema, cuando el parámetro de regularización tiende a cero; en el caso de datos con ruido se puede obtener una aproximación óptima de la solución exacta para valores del parámetro distintos de cero.

Consideremos el siguiente problema $Ax = y$ donde A es un operador compacto entre los espacios de Hilbert X y Y sobre el campo de los complejos o reales. En la práctica y nunca es conocido exactamente, supongamos que $y^\delta \in Y$ para un $\delta > 0$ con $\|y - y^\delta\| \leq \delta$, el objetivo es resolver

$$Ax^\delta = y^\delta,$$

que en general no tiene solución debido a que no necesariamente se tiene que y^δ está en el rango de A ($R(A)$), por lo tanto lo mejor que podemos hacer es determinar una aproximación $x^\delta \in X$ a la exacta x , también se requiere que la aproximada x^δ dependa continuamente de los datos y^δ . En otras palabras queremos construir una aproximación acotada adecuada $R : Y \rightarrow X$ del operador inverso no acotado $A^{-1} : R(A) \rightarrow X$.

Definición 1.1.1. Una estrategia de regularización es una familia de operadores lineales acotados

$$R_\alpha : Y \rightarrow X, \quad \alpha > 0,$$

tal que

$$\lim_{\alpha \rightarrow 0} R_\alpha Ax = 0 \text{ para toda } x \in X, \quad (1.2)$$

aquí el operador $R_\alpha A$ converge puntualmente a la identidad.

Teorema 1.1.2. Sea R_α una estrategia de regularización para un operador compacto $A : X \rightarrow Y$ donde $\dim(X) = \infty$. Entonces tenemos

- Los operadores R_α no son uniformemente acotados; esto es, existe una sucesión (α_j) con $\|R_{\alpha_j}\| \rightarrow \infty$ cuando $j \rightarrow \infty$.
- La sucesión $(R_\alpha Ax)$ no converge uniformemente sobre subconjuntos acotados de X ; es decir, no existe convergencia de $R_\alpha A$ a la identidad I .

Ejemplo 4.

Veamos una solución aproximada del **ejemplo 1** para la ecuación (1.1) por cociente de diferencias centradas. Sea $\alpha = h$

$$(R_h g)(x) := \frac{g(x+h) - g(x-h)}{2h} \quad x \in [0, 1],$$

para $h > 0$. Para hacer que $(R_h g)(x)$ esté bien definido para x cerca de las fronteras, supongamos por simplicidad que g es periódica con periodo 1. Al hacer una expansión de Taylor para g se tiene

1. Preliminares

$$\|g' - R_h g\|_\infty \leq \frac{h}{2} \|g''\|_\infty,$$

$$\|g' - R_h g\|_\infty \leq \frac{h^2}{6} \|g'''\|_\infty,$$

si $g \in C^2([0, 1])$ o $g \in C^3([0, 1])$ respectivamente. Para el dato ruidoso g^δ el error total puede ser estimado por

$$\begin{aligned} \|g' - R_h g^\delta\|_\infty &\leq \|g' - R_h g\|_\infty + \|R_h g - R_h g^\delta\|_\infty \\ &\leq \frac{h}{2} \|g''\|_\infty + \frac{\delta}{h}, \end{aligned}$$

si $g \in C^2([0, 1])$. En este estimativo tenemos el error total dividido en un error aproximado $\|g' - R_h g\|_\infty$, el cual tiende a 0 cuando $h \rightarrow 0$, y el error en el dato ruidoso $\|R_h g - R_h g^\delta\|_\infty$, el cual se va a infinito cuando $h \rightarrow 0$. Para dar una buena aproximación tenemos que balancear estos dos términos de error por una buena elección de la discretización del parámetro h . El mínimo del lado derecho para el error estimado es obtenido en $h = (2/\|g''\|)^{1/2} \delta^{1/2}$. Con esta elección de h el error total es de orden

$$\|g' - R_h g^\delta\|_\infty = O(\delta^{1/2}). \quad (1.3)$$

Si $g \in C^3([0, 1])$, haciendo un cálculo similar tenemos que $h = (3/\|g'''\|)^{1/3} \delta^{1/3}$ y tenemos una mejor razón de convergencia

$$\|g' - R_h g^\delta\|_\infty = O(\delta^{2/3}). \quad (1.4)$$

Mas suavidad de g no mejora el orden de $\delta^{2/3}$ para el cociente de diferencias centradas R_h . Esto demuestra que para orden más alto en el esquema de diferencias la razón de convergencia es más pequeño que $O(\delta)$.

La razón de convergencia (1.3) refleja el hecho de que la estabilidad puede ser restaurada en el problema (1.1) si es dada la información

$$\|g''\|_\infty \leq E.$$

En otras palabras, la restricción del operador diferencial en el conjunto $W_E := \{g : \|g''\| \leq E\}$ es continuo con respecto a la norma del máximo. Esto se sigue del siguiente estimativo

$$\begin{aligned} \|g'_1 - g'_2\|_\infty &\leq \|(\frac{d}{dx} - R_h)(g_1 - g_2)\|_\infty + \|R_h(g_1 - g_2)\|_\infty \\ &\leq \frac{h}{2} \|g''_1 - g''_2\|_\infty + \frac{\|g_1 - g_2\|_\infty}{h} \\ &\leq 2\sqrt{E\|g_1 - g_2\|_\infty}, \end{aligned}$$

el cual se tiene para $g_1, g_2 \in W_E$ con la elección $h = \sqrt{\|g_1 - g_2\|_\infty / E}$.

1. Preliminares

1.2. Molificación para la diferenciación numérica

Como definimos en la sección anterior una de las técnicas conocidas para regularizar un problema inverso mal puesto es la molificación, en esta sección mostraremos cómo funciona la molificación para aproximar derivadas de funciones que vienen con ruido, debido a la práctica ó datos experimentales.

Supongamos que tenemos una función con ruido, $f^\epsilon = f + \epsilon$, nosotros queremos estimar $J_\delta f^\epsilon$ tal que $\|f - J_\delta f^\epsilon\|$ sea pequeño, pero al mismo tiempo que sea suave. Sea $\delta > 0$, $p > 0$, y

$$A_p = \left(\int_{-p}^p e^{-s^2} ds \right)^{-1}.$$

La δ -**molificación** de una función integrable se basa sobre la convolución con el kernel Gaussiano

$$\rho_{\delta,p}(x) = \begin{cases} A_p \delta^{-1} e^{-(x^2/\delta^2)} & \text{si } |x| \leq p\delta, \\ 0 & \text{si } |x| > p\delta. \end{cases}$$

El δ -molificado, $\rho_{\delta,p}$, es no negativo en $C^\infty(I)$ donde $I = (-p\delta, p\delta)$ y satisface

$$\int_{-p\delta}^{p\delta} \rho_{\delta,p}(x) dx = 1.$$

Para $f(x) \in L^1_{loc}(\mathbb{R})$, definimos esta δ -molificación por el operador convolución

$$J_\delta f(x) = (\rho_\delta * f)(x) = \int_{-\infty}^{\infty} \rho_\delta(x-s) f(s) ds = \int_{x-p\delta}^{x+p\delta} \rho_\delta(x-s) f(s) ds.$$

Sea $K = \{x_j : j \in \mathbb{Z}\} \subset \mathbb{R}$, que satisface $x_{j+1} - x_j > d > 0$, $j \in \mathbb{Z}$ y $G = \{g_j\}_{j \in \mathbb{Z}}$ una función discreta definida sobre K , además sea $s_j = \frac{1}{2}(x_{j+1} + x_j)$, $j \in \mathbb{Z}$, entonces la δ -molificación de G para el caso discreto se define como

$$J_\delta G(x) = \sum_{j=-\infty}^{\infty} \left(\int_{s_{j-1}}^{s_j} \rho_\delta(x-s) ds \right) g_j,$$

y note que

$$\sum_{j=-\infty}^{\infty} \left(\int_{s_{j-1}}^{s_j} \rho_\delta(x-s) ds \right) = \int_{x-p\delta}^{x+p\delta} \rho_\delta(s) ds = 1.$$

1.2.1. Aproximación de derivadas

Supongamos que damos una función con ruido f^ϵ y queremos estimar su derivada f' , haciendo uso de las propiedades del operador convolución, nosotros podemos proceder de diferentes maneras:

1. Preliminares

- Usar la definición directamente $(\rho_\delta * f)'$.
- Primero calcular f' y después hacer la convolución con ρ'_δ .

Las propiedades principales que relacionan la δ -molificación con la diferenciación de funciones con ruido son establecidas en los siguientes resultados, [9].

Teorema 1.2.1.

1. (Consistencia) Si $f'(x) \in C^{0,1}(\mathbb{R})$, entonces existe una constante C , independiente de δ , tal que

$$\|(J_\delta f)' - f\|_\infty \leq C\delta.$$

2. (Estabilidad) Si $f(x), f^\varepsilon(x) \in C^0(\mathbb{R})$ y $\|f - f^\varepsilon\|_\infty \leq \epsilon$, existe una constante C , independiente de δ , tal que

$$\|(J_\delta f)' - (J_\delta f^\varepsilon)'\|_\infty \leq C \frac{\epsilon}{\delta}.$$

3. (Convergencia) $\|(J_\delta f^\varepsilon)' - f'\|_\infty \leq C \left(\delta + \frac{\epsilon}{\delta} \right).$

Ahora veamos los resultados para la molificación discreta. Sea $\Delta x = x_{j+1} - x_j$ para $j \in \mathbb{Z}$.

Teorema 1.2.2.

1. (Consistencia) Sea $g'(x) \in C^{0,1}(\mathbb{R}^1)$, y $G = \{g_j = g(x_j) : j \in \mathbb{Z}\}$ la versión discreta de g . Entonces existe una constante C , independiente de δ , tal que

$$\|(J_\delta G)' - g'\|_\infty \leq C \left(\delta + \frac{\Delta x}{\delta} \right).$$

2. (Estabilidad) Para las funciones discretas G y $G^\varepsilon = \{g_j^\varepsilon : j \in \mathbb{Z}\}$, definidas en $\mathbb{R}$ con $\|G - G^\varepsilon\|_\infty \leq \epsilon$, se tiene que

$$\|(J_\delta G)' - (J_\delta G^\varepsilon)'\|_\infty \leq 4A_p \frac{\epsilon}{\delta}.$$

3. (Convergencia) Sea $g'(x) \in C^{0,1}(\mathbb{R}^1)$ y $G = \{g_j = g(x_j) : j \in \mathbb{Z}\}$ la versión discreta de g , y $G^\varepsilon = \{g_j^\varepsilon : j \in \mathbb{Z}\}$ es la versión discreta perturbada de g que satisface

$$\|G - G^\varepsilon\|_\infty \leq \varepsilon.$$

Entonces existe una constante C , independiente de δ , tal que

$$\|(J_\delta G^\varepsilon)' - (J_\delta g)'\|_\infty \leq \frac{C}{\delta} (\epsilon + \Delta x),$$

y

$$\|(J_\delta G^\varepsilon)' - g'\|_\infty \leq C \left(\delta + \frac{\epsilon}{\delta} + \frac{\Delta x}{\delta} \right).$$

1. Preliminares

1.2.2. Diferenciación numérica por molificación discreta

La diferenciación numérica de datos inexactos es un problema mal formulado, y el método que se presenta aquí se sigue de la reconstrucción estable de la derivada de una función que es aproximadamente conocida en un conjunto discreto de puntos.

Sea $G^\varepsilon = \{g_j^\varepsilon : j \in \mathbb{Z}\}$ la función discreta perturbada de g . Para reconstruir la derivada g' de datos con ruido utilizamos $(d/dx)\rho_\delta$ y la convolución con la función dada. Los cálculos son arreglados con la aproximación de diferencias centradas de la derivada molificada $(d/dx)J_\delta G^\varepsilon$, donde el operador de diferencia central (D_0) está dado por

$$D_0 f(x) = \frac{f(x + \Delta x) - f(x - \Delta x)}{2\Delta x}.$$

Teorema 1.2.3.

Suponga que $g'(x) \in C^{0,1}(\mathbb{R})$ y sean $G = \{g_j = g(x_j) : j \in \mathbb{Z}\}$ la versión discreta de g , y $G^\varepsilon = \{g_j^\varepsilon : j \in \mathbb{Z}\}$ la versión discreta perturbada de g que satisface $\|G - G^\varepsilon\|_\infty \leq \varepsilon$. Entonces

$$\|D_0(J_\delta G^\varepsilon) - (J_\delta g)'\|_\infty \leq \frac{C}{\delta} (\varepsilon + \Delta x) + C_\delta (\Delta x)^2$$

y

$$\|D_0(J_\delta G^\varepsilon) - g'\|_\infty \leq C \left(\delta + \frac{\varepsilon}{\delta} + \frac{\Delta x}{\delta} \right) + C_\delta (\Delta x)^2,$$

en donde el operador diferenciación D_0^δ se define por la siguiente regla: $D_0^\delta = D_0(J_\delta G)(x)|_{\mathbb{R}}$.

Teorema 1.2.4.

El operador D_0 satisface

$$\|D_0^\delta(G)\|_\infty \leq 4 \frac{A_p}{\delta} \|G\|_\infty.$$

Ver [9].

1.2.3. Ejemplos

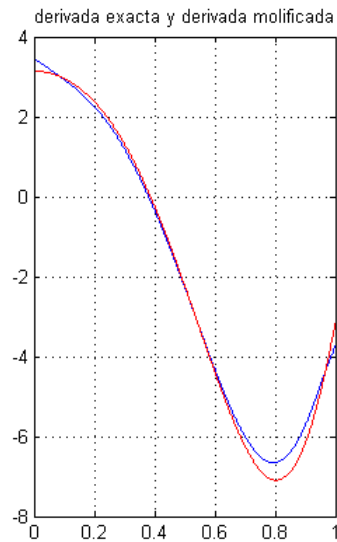
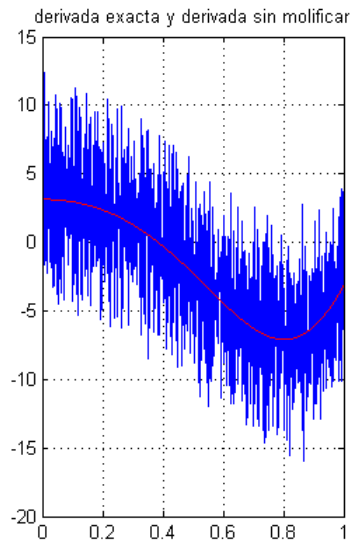
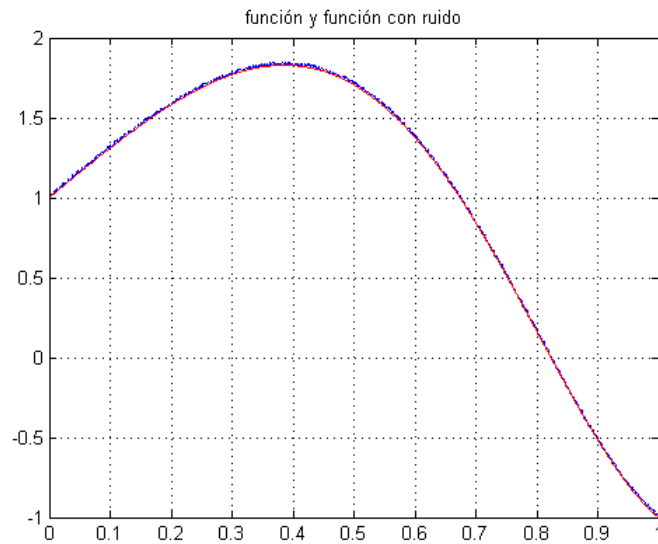
Ejemplo 5.

Sean

$$f(x) = \sin(\pi x) + \cos(-\pi x^2) \quad y \quad f'(x) = \pi \cos(\pi x) + 2\pi x \sin(-\pi x^2),$$

las función exactas y $\varepsilon = 0,02$ el ruido para la función f .

1. Preliminares



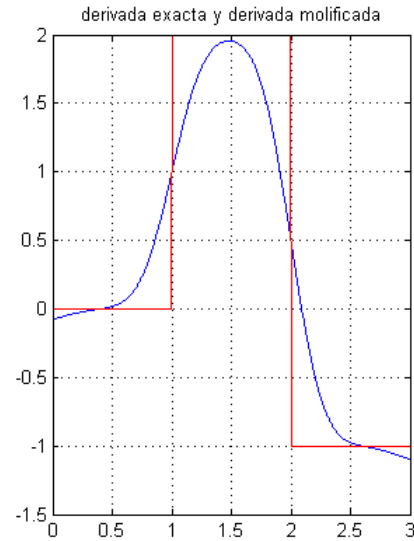
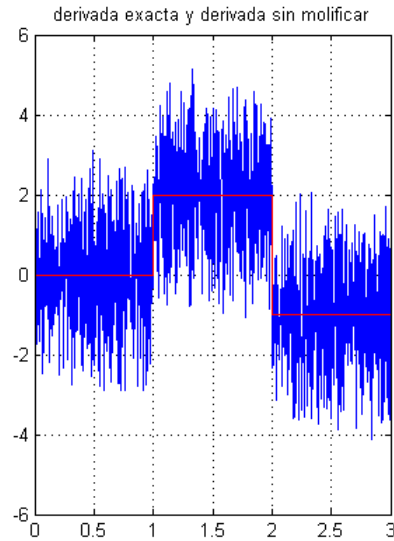
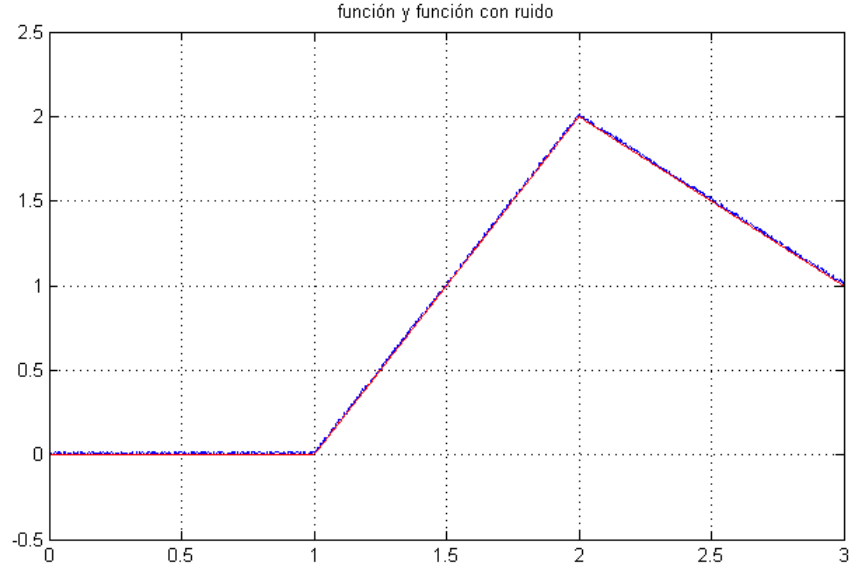
Ejemplo 2

Sean

$$f(x) = \begin{cases} 0 & \text{si } 0 \leq x < 1 \\ 2x - 2 & \text{si } 1 \leq x < 2 \\ -x + 4 & \text{si } 2 \leq x \leq 3 \end{cases} \quad y \quad f'(x) = \begin{cases} 0 & \text{si } 0 \leq x < 1 \\ 2 & \text{si } 1 \leq x < 2 \\ -1 & \text{si } 2 \leq x \leq 3 \end{cases},$$

las funciones exactas y $\epsilon = 0,02$ el ruido para la función f .

1. Preliminares



1.3. Convergencia débil

Definición 1.3.1. La clausura convexa de un conjunto A en un espacio vectorial X es el conjunto de todas las combinaciones convexas de elementos de A . Esto es, el conjunto de todas las sumas $t_1x_1 + t_2x_2 + \dots + t_nx_n$ en las cuales $x_i \in A, t_i \geq 0, \sum t_i = 1$; n es arbitrario.

Proposición 1.3.2. Sean X, Y espacios normados. Si Y es un espacio de Banach, entonces $B(X, Y)$ es un espacio de Banach.

1. Preliminares

De esta proposición se deduce que $X' = B(X, \mathbb{R})$ siempre es un espacio de Banach (aún cuando X no lo sea). Más aún si $f \in X'$ y $u \in X$ se suele denotar

$$\langle f, u \rangle = f(u)$$

y se denomina producto de dualidad de f con u .

Definición 1.3.3. Sea X un espacio de Banach y sea J la inyección canónica de X en X'' (bidual de X). Se dice que X es reflexivo si $J(X) = X''$.

Definición 1.3.4. Sea X un espacio normado. Una sucesión $\{u_n\} \subset X$ se dice que converge débilmente a un punto $u \in X$ y se escribe

$$u_n \rightharpoonup u,$$

si y sólo si

$$\forall f \in X' \text{ se tiene } \langle f, u_n \rangle \rightarrow \langle f, u \rangle.$$

Veamos unas propiedades fundamentales de la convergencia débil sus demostraciones se pueden ver en [18].

Propiedades:

Una sucesión $\{u_n\}$ en un espacio de Banach X sobre el campo $\mathbb{R}$, entonces:

1. La convergencia fuerte $u_n \rightarrow u$ cuando $n \rightarrow \infty$ implica la convergencia débil $u_n \rightharpoonup u$ cuando $n \rightarrow \infty$.
2. Si $\dim X < \infty$, entonces la convergencia débil $u_n \rightharpoonup u$ cuando $n \rightarrow \infty$ implica la convergencia fuerte $u_n \rightarrow u$ cuando $n \rightarrow \infty$.
3. Si $u_n \rightharpoonup u$ cuando $n \rightarrow \infty$, entonces $\{u_n\}$ es acotada y $\|u\| \leq \limsup \|u_n\|$ cuando $n \rightarrow \infty$. (Teorema de Banach y Steinhaus 1927).
4. Si $u_n \rightharpoonup u$ cuando $n \rightarrow \infty$, entonces existe una sucesión $\{v_n\}$ en la clausura convexa cerrada de $\{u_n\}$ tal que $v_n \rightarrow u$ cuando $n \rightarrow \infty$. (Teorema de Mazur 1933).
5. Si $u_n \rightharpoonup u$ cuando $n \rightarrow \infty$, entonces u está en la clausura convexa cerrada de $\{u_n\}$.
6. Si $\{u_n\}$ es acotado en X y si existe un $u \in X$ y un conjunto denso D en X' tal que $\langle f, u_n \rangle \rightarrow \langle f, u \rangle$ cuando $n \rightarrow \infty$ para todo $f \in D$, entonces $u_n \rightharpoonup u$ cuando $n \rightarrow \infty$.
7. Si X es reflexivo y la sucesión real $\{\langle v, u_n \rangle\}$ converge para todo $v \in X'$, entonces existe un $u \in X$ tal que $u_n \rightharpoonup u$ cuando $n \rightarrow \infty$.
8. Si $\{u_n\}$ es una sucesión acotada en el B-espacio reflexivo X , entonces $\{u_n\}$ tiene una subsucesión que converge débilmente. Si además cada subsucesión de $\{u_n\}$ converge débilmente a u , entonces $u_n \rightharpoonup u$ cuando $n \rightarrow \infty$.

1. Preliminares

9. Se sigue de

$$u_n \rightharpoonup u \text{ en } X \text{ cuando } n \rightarrow \infty.$$

$$f_n \rightarrow f \text{ en } X' \text{ cuando } n \rightarrow \infty,$$

$$\text{que } \langle f_n, u_n \rangle \rightarrow \langle f, u \rangle \text{ cuando } n \rightarrow \infty.$$

10. Si X es reflexivo, entonces se sigue de

$$u_n \rightharpoonup u \text{ en } X \text{ cuando } n \rightarrow \infty.$$

$$f_n \rightharpoonup f \text{ en } X' \text{ cuando } n \rightarrow \infty,$$

$$\text{que } \langle f_n, u_n \rangle \rightarrow \langle f, u \rangle \text{ cuando } n \rightarrow \infty.$$

1.4. Métodos iterativos

A. Método del descenso rápido

Consideremos un sistema de ecuaciones lineales de la forma

$$\mathbf{A}x = b, \tag{1.5}$$

donde $\mathbf{A}$ es de la forma $A + \alpha G$, A y G son matrices de orden N , simétricas con A semidefinida positiva, G es definida positiva, y $\alpha \geq 0$ es un parámetro dado. Suponga que los autovalores de A y G son $\{\mu_k\}_{k=1}^N$ y $\{\lambda_k\}_{k=1}^N$ respectivamente con

$$0 \leq \mu_k \leq \mu_{k+1} \text{ para } 1 \leq k \leq N,$$

$$0 < \lambda_k \leq \lambda_{k+1} \text{ para } 1 \leq k \leq N.$$

Se define el número de condición de la matriz A como $\kappa(A) := \frac{\mu_N}{\mu_1}$ y el número de condición de G como $\kappa(G) \leq C$, es decir, acotado independientemente de N . De hecho si indicamos la dependencia de los autovalores sobre el tamaño de la matriz, esto es $\mu_k = \mu_k^{(N)}$, entonces frecuentemente se tiene:

$$\lim_{n \rightarrow \infty} \mu_1^{(n)} = 0, \mu_k^{(N)} \text{ es acotado cuando } N \rightarrow \infty \text{ luego } \lim_{N \rightarrow \infty} \kappa(A) = +\infty.$$

Observación:

En problemas directos (de ecuaciones diferenciales parciales) es común encontrar que la matriz que resulta de la discretización (por elementos finitos o diferencias finitas) tiene un número de condición grande, esto es debido a que el máximo autovalor es muy grande en comparación con el más pequeño. Además la matriz usualmente es esparcida y muy grande.

1. Preliminares

Veamos, para $\alpha > 0$, que $\kappa(A + \alpha G) \leq \frac{\mu_N + \alpha \lambda_N}{\mu_1 + \alpha \lambda_1} < \kappa(A)$ si $\kappa(A) > \kappa(G)$.

El descenso rápido es basado en la caracterización de la solución de (1.5) como el minimizador del funcional

$$f(x) = (\mathbf{A}x, x) - 2(x, b), \quad x \in \mathbb{R}^N,$$

dados una matriz $\mathbf{A}$ y un vector b donde $(\cdot, \cdot)$ denota el producto interior euclidiano sobre $\mathbb{R}^N$.

Entonces no es difícil probar el siguiente resultado:

Proposición 1.4.1. *Si $\mathbf{A}$ es una matriz simétrica de orden N definida positiva, entonces*

$$\mathbf{A}u = b \iff f(u) = \min_{x \in \mathbb{R}^N} f(x).$$

Supongamos que tenemos alguna aproximación de $u = \mathbf{A}^{-1}b$, digamos $y \approx u$. Ahora consideremos una nueva aproximación de la forma

$$w = y + tv,$$

donde $t \in \mathbb{R}$ y $v \in \mathbb{R}^N$. Necesitamos elegir la dirección de v tal que el funcional tenga la razón de decrecimiento (cerca de y) mayor. Del cálculo multivariado conocemos que $\nabla f(y)$ da la dirección de mayor razón de cambio de f en y y por lo tanto tenemos que $v = -\nabla f(y)$. Con un simple cálculo tenemos

$$\nabla f(y) = \mathbf{A}y - b.$$

Con la elección de la dirección $v = -\nabla f(y)$ elegimos el tamaño del paso $t = t^*$ minimizando $F(t) := f(y + tv)$. Esto es $F'(t^*) = 0$ y por tanto

$$t^* = \frac{\|v\|^2}{(\mathbf{A}v, v)}.$$

Consecuentemente, la expectativa es que el nuevo vector, $w = y + t^*v$ con $v = b - \mathbf{A}y$, es la mejor aproximación de u para y .

Algoritmo de Descenso Rápido: Paso inicial $y^{(0)} \in \mathbb{R}^N$

Para $k \geq 0$:

poner $r^{(k)} = b - \mathbf{A}y^{(k)}$ y PARE si $r^{(k)} = 0$

calcule $t_k = \frac{\|v\|^2}{(\mathbf{A}v, v)}$ y actualizar $y^{(k+1)} = y^{(k)} + t_k r^{(k)}$.

1. Preliminares

Observación:

Si $\mathbf{A}$ es una matriz cuadrada simétrica definida positiva entonces podemos definir un producto interior sobre $\mathbb{R}^N$ por

$$[x, y] = (\mathbf{A}x, y).$$

No es difícil mostrar que para $\alpha, \beta \in \mathbb{R}$ y $x, y, v \in \mathbb{R}^N$

- $[x, x] \geq 0$ y $[x, x] = 0 \iff x = 0$.
- $[x, y] = [y, x]$.
- $[\alpha x + \beta v, y] = \alpha[x, y] + \beta[v, y]$.

La correspondiente norma es dada por

$$\|x\|_{\mathbf{A}} = \sqrt{(\mathbf{A}x, x)} \quad , \quad x \in \mathbb{R}^N.$$

Teorema 1.4.2. Para algún $k \geq 0$, el error $\varepsilon^{(k)} := u - y^{(k)}$ satisface

$$\|\varepsilon^{(k)}\|_{\mathbf{A}} \leq \left[\frac{\kappa(\mathbf{A}) - 1}{\kappa(\mathbf{A}) + 1} \right]^k \|\varepsilon^{(0)}\|_{\mathbf{A}}.$$

La prueba de este teorema se puede ver en [1].

B. Método del Gradiente Conjugado

Una desventaja del método del descenso rápido es su carácter local. Es decir, que para hallar la nueva dirección descendente se basa solamente en la ubicación actual. Consideremos entonces la dirección del paso k , $r^{(k)} = b - \mathbf{A}y^{(k)}$.

Si $y^{(0)} = 0$, tenemos $r^{(0)} = b$, y por lo tanto

$$y^{(1)} = 0 + \tau_1 b \in \text{gen}\{b\}, \quad r^{(1)} = b - \mathbf{A}y^{(1)} = b - \tau_1 \mathbf{A}b,$$

y de aquí

$$y^{(2)} = y^{(1)} + \tau_2 r^{(1)} = \tau_1 b + \tau_2 (b - \tau_1 \mathbf{A}b) \in \text{gen}\{b, \mathbf{A}b\},$$

y así sucesivamente

$$y^{(k)} \in \text{gen}\{b, \mathbf{A}b, \dots, \mathbf{A}^{k-1}b\} = S_k.$$

El espacio S_k es llamado el k -ésimo subespacio de Krylov correspondiente a b . Note también que

$$\text{gen}\{\mathbf{A}x, \mathbf{A}^2x, \dots, \mathbf{A}^kx\} = S_k, \quad \text{donde } x = \mathbf{A}^{-1}b.$$

En el método del gradiente conjugado (GC) se usa el dato anterior para determinar la dirección descendente. Este método aplicado al sistema $\mathbf{A}x = b$ puede ser definido como sigue.

1. Preliminares

Para determinar $w^{(k)}$ en la k -ésima (GC) iteración, hacemos uso del método de Galerkin con el subespacio S_k , de la siguiente manera

$$(\mathbf{A}w^{(k)} - b, v) = 0, \quad v \in S_k,$$

o equivalentemente

$$[w^{(k)} - x, v] = 0, \quad v \in S_k,$$

por lo tanto la aproximación k -ésima (GC) satisface

$$[w^{(k)}, v] = (b, v), \quad v \in S_k. \quad (1.6)$$

Esto significa que $w^{(k)}$ es la proyección de x sobre S_k respecto al producto interior $[\cdot, \cdot]$. No es necesario tener un orden para conocer x y a su vez calcular esta proyección. En efecto,

$$\|x - w^{(k)}\|_{\mathbf{A}} \leq \|x - v\|_{\mathbf{A}}, \quad v \in S_k.$$

En particular,

$$\|x - w^{(k)}\|_{\mathbf{A}} \leq \|x - y^{(k)}\|_{\mathbf{A}},$$

donde $\{y^{(k)}\}$ es la sucesión de descenso empinado.

No es del todo obvio cómo puede ser calculado w^k por un algoritmo relativamente simple. No probaremos esto aquí simplemente identificamos el algoritmo. Ver [1].

Algoritmo del Gradiente Conjugado:

entrar $p^{(0)} = r^{(0)} = b$, $w^{(0)} = 0$

calcule para $k \geq 0$:

$$\mathbf{A}p^{(k)}$$

$$w^{(k+1)} = w^{(k)} + \alpha_k p^{(k)}, \quad \text{donde } \alpha_k = \frac{(r^{(k)}, p^{(k)})}{(\mathbf{A}p^{(k)}, p^{(k)})}$$

$$r^{(k+1)} = r^{(k)} - \alpha_k \mathbf{A}p^{(k)}$$

$$p^{(k+1)} = r^{(k+1)} - \beta_k p^{(k)}, \quad \text{donde } \beta_k = \frac{(\mathbf{A}p^{(k)}, r^{(k+1)})}{(\mathbf{A}p^{(k)}, p^{(k)})}.$$

El próximo teorema da un resultado de convergencia para el método del gradiente conjugado.

Teorema 1.4.3. *Para algún $k \geq 0$,*

$$\|\varepsilon^{(k)}\|_{\mathbf{A}} \leq 2 \left[\frac{\kappa^{1/2}(\mathbf{A}) - 1}{\kappa^{1/2}(\mathbf{A}) + 1} \right]^k \|\varepsilon^{(0)}\|_{\mathbf{A}},$$

donde $\varepsilon^{(k)} = x - w^{(k)}$. Además la sucesión $\{w^{(k)}\}$ del GC satisface

$$\|x - w^{(k)}\|_{\mathbf{A}} = \min_{w \in S_k} \|x - w\|_{\mathbf{A}}.$$

Estos métodos se pueden encontrar en [1].

Capítulo 2

Ecuación del calor

2.1. Derivación de las ecuaciones reacción-difusión

Las ecuaciones de reacción-difusión implican dos procesos diferentes: difusión (movimiento local) y reacción (crecimiento, interacciones, cambios de estado).

La difusión es un fenómeno el cual un grupo de partículas se mueven como grupo de acuerdo a la trayectoria irregular de cada una de las partículas. El principio de conservación es una de las teorías físicas básicas en la formulación de ecuaciones. Cuando el problema en cuestión consiste en un proceso de reacción acompañada de la difusión, este principio conduce a un conjunto de ecuaciones en derivadas parciales de las incógnitas del sistema. Estas cantidades podrán ser: concentraciones de masas en los procesos de reacción química, la temperatura en la conducción del calor, flujo de neutrones en los reactores nucleares, la densidad de población en la dinámica de la población y muchos otros. Para dar una descripción de la deducción de las ecuaciones que se presentan consideremos $u(x, t)$, como la función de densidad, en el tiempo t y la posición x en un medio de difusión Ω en $\mathbb{R}^n$. El principio de conservación establece: *Para cualquier subdominio R en Ω , con frontera cerrada S , la razón de cambio de la densidad de masa en R es igual a menos la razón del flujo a través de la frontera S .* Sea F el campo vectorial denominado flujo y sea ν el vector normal exterior a S . Supongamos que u , F son continuas en x y que F tiene derivadas parciales continuas con respecto a las componentes de x y t además que u tiene una derivada continua en t . Entonces la relación de equilibrio se puede expresar como

$$\frac{d}{dt} \int_R a_0 u dx = - \oint_S F \cdot \nu ds, \quad (2.1)$$

donde a_0 es una constante de proporcionalidad, la cual depende del tipo de problema que se está estudiando. Por ejemplo, en procesos de reacción química, a_0 representa el número de Lewis, y en problemas de conducción de calor es el producto entre la densidad y el calor específico. El signo negativo en la integral de superficie representa el flujo de densidad en la región R a través de la frontera S . Del teorema de la divergencia tenemos que

2. Ecuación del calor

$$\oint_S F \cdot \nu ds = \int_R \nabla \cdot F dx,$$

donde ∇ es el operador gradiente en x , y la ecuación (2.1) se reduce a

$$\int_R (a_0 u_t + \nabla \cdot F) dx = 0.$$

Suponiendo que u_t y $\nabla \cdot F$ son continuas en Ω y la arbitrariedad del subdominio R implica que

$$a_0 u_t + \nabla \cdot F = 0, \text{ en } \Omega. \quad (2.2)$$

Para relacionar el flujo de difusión F para la función de densidad u necesitamos de un principio físico. Este principio tiene diferentes nombres en diferentes contextos. Se llama la ley de Fick en los procesos de reacción química, la ley de Fourier en los problemas de la conducción del calor y la ley de Darcy en las ecuaciones de medio poroso. En cada caso, la ley establece, el flujo es proporcional al gradiente negativo de la densidad

$$F = -a^* \nabla u,$$

donde a^* es una función estrictamente positiva en Ω . Sustituyendo esta relación en (2.2) tenemos la ecuación

$$u_t = \nabla \cdot (a \nabla u), \quad (2.3)$$

donde $a = a^*/a_0$. La función a es llamada el coeficiente de difusión en procesos de difusión químicos ó el término de difusividad en problemas de conducción de calor. El término $\nabla \cdot (a \nabla u)$ representa la razón de cambio debido a la difusión.

Cuando tenemos factores externos como la fuente en nuestro fenómeno físico, el principio de conservación establece que

$$\frac{d}{dt} \int_R a_0 u dx = - \oint_S F \cdot \nu ds + \int_R q dx.$$

En problemas de conducción de calor podemos tener fuentes q externas que generen más calor o sumideros donde se presenta pérdida de calor, en este caso, el término reacción q es la densidad por unidad de volumen por unidad de tiempo formado a través del proceso de reacción ó interacción. Haciendo un proceso análogo a la deducción de la ecuación tenemos que

$$u_t = \nabla \cdot (a \nabla u) + q. \quad (2.4)$$

Cuando q es una función determinada por la ecuación estándar (2.4) esta es la ecuación estándar difusión lineal o ecuación del calor. En muchos problemas de tipo reacción-difusión, q depende de la función u densidad y posiblemente en (x, t) de forma explícita. Escribir $q = F(x, t, u)$ en (2.4) conduce a la ecuación de reacción-difusión

2. Ecuación del calor

$$u_t - \nabla \cdot (a \nabla u) = F(x, t, u). \quad (2.5)$$

La ecuación (2.5) en el campo de la ciencia aplicada es llamada ecuación reacción difusión no estacionaria. En la literatura matemática a menudo son conocidas como ecuaciones parabólicas semilineales .

Más información sobre las ecuaciones reacción-difusión se puede ver en [11].

2.2. Aplicaciones

En el contexto de la conducción de calor y la difusión cuando u representa la temperatura y la concentración, la función $F(x, t, u)$ desconocida se interpreta como una fuente de calor y material. Por otro lado en aplicaciones de productos químicos o bioquímicos se puede interpretar como un término de reacción. Veamos unas aplicaciones.

1. Modelo lineal de Fischer del avance de los genes ventajosos

Considere la población de una especie dada, distribuida uniformemente a lo largo de un hábitat lineal. Supongamos que el tamaño del hábitat es grande en comparación con las distancias que separan los lugares de la descendencia de aquellos que son de sus padres. Supongamos además que en algún punto del hábitat una mutación ventajosa ocurre (es decir, una mutación que es de alguna manera ventajosa para la supervivencia de algún miembro de la población).

Esta mutación se difunde, primero en la vecindad de la ocurrencia de la mutación y después en los alrededores de la población.

Sea $p = p(x, t)$ la concentración de los miembros de la población con el gen mutante y sea $q = q(x, t)$ la concentración de los miembros de la población cuyos descendientes tienen el gen mutante (x es la posición dentro el hábitat). Podemos asumir que $q = 1 - p$. Denotemos por α la intensidad de la selección en favor del gen mutante, el cual tomamos como independiente de p . Para α suficientemente pequeña, la concentración p varía continuamente con el tiempo de generación en generación. Suponga que la razón por generación a la cual los miembros de la población con el gen mutante se difunde dentro de la población total está dada por $-k \frac{\partial p}{\partial x}$, donde $k > 0$ es una constante de difusión (independiente de x y p). Bajo todas estas suposiciones se puede mostrar que la concentración del gen mutante satisface

$$\frac{\partial p}{\partial t} = \frac{\partial^2 p}{\partial x^2} + \alpha p(1 - p).$$

Ver [13].

2. Ecuación del calor

2. Genética poblacional

Consideremos una población de individuos en donde cada uno porta dos genes, digamos B y A, así, la población se divide en tres genotipos BB, BA y AA. Sea Ω el hábitat donde vive la población y sean $u_i(x, t)$, $i = 1, 2, 3$, las densidades de población de BB, BA y AA respectivamente. Supongamos que la población se cruza de manera aleatoria con una tasa de nacimiento r y que la población se mueve con una constante de difusión a . También suponemos que las tasas de muerte dependen de los genotipos y los denotamos respectivamente con τ_i , $i = 1, 2, 3$. Bajo estas condiciones las densidades u_i satisfacen el sistema

$$\frac{\partial u_i}{\partial t} = \nabla(a_i \nabla u_i) + f_i(x, t, u_1, u_2, u_3), i = 1, 2, 3.$$

Ver [13].

3. Dispersión de mamíferos

Modelos espaciales para la dispersión de mamíferos en ecología vienen desde Skellam [20], quien modeló la expansión de la población de muskrats (roedores acuáticos) en Europa. Skellam encontró una relación lineal entre la raíz cuadrada del área habitada por los muskrats y el tiempo, usando un modelo de dos dimensiones con dispersión aleatoria y crecimiento exponencial de la población en coordenadas polares. Suponiendo que la dispersión se da por igual en todas direcciones, es decir, es isotrópica, y que el crecimiento de la población es proporcional a la densidad poblacionales, el modelo es:

$$\frac{\partial S}{\partial t} = \frac{a}{r} \frac{\partial}{\partial r} \left(r \frac{\partial S}{\partial r} \right) + \alpha S,$$

donde a es la constante de difusión y α la tasa neta de crecimiento. Ver [13].

4. Modelo de Maltus

En la predicción de la mayoría de la población, el modelo de Malthus sólo se puede utilizar en un corto período de tiempo. Debido a la limitación de los recursos, la tasa de crecimiento de la población se reducirá y la población pueden saturar a un nivel máximo. Por lo tanto, es natural que se utilice una tasa de crecimiento que depende de la densidad por habitante, la más simple que se utiliza sobre todo, es la tasa de crecimiento logístico. Junto con la difusión de la especie, se obtiene en la siguiente ecuación logística difusiva

$$\frac{\partial P}{\partial t} = D \Delta P + aP \left(1 - \frac{P}{N} \right),$$

2. Ecuación del calor

donde $a > 0$ es la máxima razón de crecimiento por capital y $N > 0$ es la capacidad de carga. La ecuación logística difusiva es la más simple ecuación no lineal de reacción-difusión, pero los métodos de análisis pueden ser utilizados en muchos otros problemas.

5. Reactores químicos y combustión

En la reacción química no isotérmica el proceso participa de múltiples especies químicas, las concentraciones químicas normalizadas y la temperatura son gobernadas por un sistema acoplado de ecuaciones de reacción-difusión. En el caso de una sola especie donde el coeficiente de difusión y la difusividad térmica son constantes las ecuaciones para la concentración u y la temperatura v son dadas por

$$\begin{aligned}u_t - D_1 \nabla^2 u &= -a_1 f(u, v), \\v_t - D_2 \nabla^2 v &= a_2 f(u, v),\end{aligned}\tag{2.6}$$

donde a_1 se conoce como el número Thiele y a_2/a_1 es la temperatura Prater. Si la reacción es irreversible la función f es dada

$$f(u, v) = u^p r(v),$$

y se conoce como la reacción de orden p . De acuerdo a la cinética de Arrhenius la función $r(v)$ es determinada por la ecuación

$$\frac{d}{dv}(\ln r) = E/Rv^2,$$

donde E es la energía de activación y R es el gas constante. Esto da la relación

$$r(v) = r_0 e^{(-E/Rv)} \equiv e^{(\gamma - \gamma/v)},$$

donde $\gamma = E/R$ y $r_0 = e^{(\gamma)}$. Así la función reacción f es

$$f(u, v) = u^p e^{(\gamma - \gamma/v)},$$

donde γ es número Arrhenius y p es el orden de la reacción. Ver [11].

6. Explosión térmica

En la derivación de las ecuaciones para el problema de reacción química, varias aproximaciones son usadas para reducir el sistema acoplado (2.6) en un problema de valor de frontera escalar. Otro modelo simple es el también llamado reacción de orden cero, el cual corresponde a $p = 0$ por la función

$$f(u, v) = u^0 e^{(\gamma - \gamma/v)}.$$

En esta situación la ecuación para la temperatura v es reducida a

$$v_t - D_2 \nabla^2 v = a_2 e^{(\gamma - \gamma/v)}.$$

Ver [11].

2. Ecuación del calor

2.3. Teoremas de existencia y unicidad de la solución

Recordemos que nuestro problema es identificar F en

$$u_t(x, t) = u_{xx}(x, t) + F(x, t, u), \quad 0 \leq x \leq 1, \quad 0 \leq t \leq 1, \quad (2.7)$$

$$\begin{aligned} u(x, 0) &= f(x) \quad 0 \leq x \leq 1, \\ u(1, t) &= g(t) \quad 0 \leq t \leq 1, \\ u(0, t) &= h(t) \quad 0 \leq t \leq 1, \end{aligned} \quad (2.8)$$

es necesario establecer las condiciones debe cumplir F para encontrar una única solución con base en las condiciones iniciales y de frontera, lo cual garantizará existencia y unicidad en la solución de nuestro problema. A continuación listaremos resultados cuyas pruebas se pueden ver en [16].

Primero supongamos que $F(x, t, u) = F(x, t)$ y definamos $\Omega = [0, 1] \times [0, 1]$.

Definición 2.3.1. La función

$$\phi(x, t) = \begin{cases} \frac{1}{(4\pi t)^{1/2}} e^{-(x^2/4t)} & \text{si } x \in \Omega, \\ 0 & \text{en otro caso} \end{cases},$$

es la solución fundamental del problema $u_t = u_{xx}$.

Donde ϕ cumple con las siguientes propiedades

- $\phi(x, t) > 0, \forall x \in [0, 1]$.
- $\int_0^1 \phi(x, t) dx = 1$ para cada $t \in (0, 1)$

Consideremos el problema de valor inicial

$$\begin{aligned} u_t - u_{xx} &= 0, \quad (x, t) \in \Omega \\ u(x, 0) &= f(x), \quad x \in [0, 1] \end{aligned}$$

y definamos la siguiente convolución

$$u(x, t) = \int_0^1 \phi(x - y, t - s) f(y) dy = \frac{1}{(4\pi t)^{1/2}} \int_0^1 e^{-(|x-y|^2/4t)} f(y) dy. \quad (2.9)$$

Teorema 2.3.2. Suponga que $f \in C([0, 1]) \cap L^\infty([0, 1])$ y definamos u como (2.9). Entonces

- $u \in C^\infty(\Omega)$,
- $u_t(x, t) - u_{xx}(x, t) = 0, (x, t) \in \Omega$,

2. Ecuación del calor

- $\lim_{\substack{(x,t) \rightarrow (x^0,0), \\ (x,t) \in \Omega}} u(x,t) = f(x^0)$ para cada punto $x^0 \in [0,1]$

Ahora consideremos el problema de valor inicial no homogéneo

$$\begin{aligned} u_t - u_{xx} &= F, & (x,t) \in \Omega \\ u(x,0) &= 0, & x \in [0,1] \end{aligned}$$

y definamos la siguiente convolución

$$u(x,t) = \int_0^t \int_0^1 \phi(x-y, t-s) F(y,s) dy ds = \int_0^t \frac{1}{(4\pi(t-s))^{1/2}} \int_0^1 e^{-(|x-y|^2/4(t-s))} F(y,s) dy ds. \quad (2.10)$$

Teorema 2.3.3. *Sea u como (2.10), $F \in C_1^2(\Omega)^1$ y de soporte compacto. Entonces*

- $u \in C_1^2(\Omega)$,
- $u_t(x,t) - u_{xx}(x,t) = F(x,t)$, $(x,t) \in \Omega$,
- $\lim_{\substack{(x,t) \rightarrow (x^0,0), \\ (x,t) \in \Omega}} u(x,t) = 0$ para cada punto $x^0 \in [0,1]$

Así, de los teoremas 2.3.2 y 2.3.3 tenemos que

$$u(x,t) = \int_0^1 \phi(x-y, t-s) f(y) dy + \int_0^t \int_0^1 \phi(x-y, t-s) F(y,s) dy ds$$

con las hipótesis sobre F y f dadas en los teoremas, u es solución de (2.7)-(2.8).

Como estamos en un dominio acotado y tenemos condiciones de frontera, la unicidad está dada por el siguiente teorema.

Teorema 2.3.4. *Existe a lo más una solución $u \in C_1^2(\Omega)$ del problema (2.7)-(2.8).*

Por último supongamos que $F := F(x,t,u)$, caso para el cual F depende de la solución. A continuación se presentarán dos teoremas que me garantizarán la existencia y unicidad del problema (2.7)-(2.8).

Teorema 2.3.5. *Sea $F(x,t,u)$ una función monotona no decreciente en u . Entonces existe a lo más una solución del problema (2.7)-(2.8).*

Teorema 2.3.6. *Sea $F(x,t,u)$ una función Lipschitz continua en u . Entonces existe a lo más una solución del problema (2.7)-(2.8).*

Las pruebas se pueden encontrar en [14].

¹ $C_1^2(\Omega) = \{F : \Omega \rightarrow \mathbb{R} : F, F_x, F_{xx}, F_t \in C(\Omega)\}$

2.4. Solución numérica

Para calcular la solución numérica de la ecuación del calor aplicaremos un método que consiste en aproximar las derivadas parciales por expresiones algebraicas llamadas diferencias finitas, veamos cómo aplicar este método para la siguiente ecuación.

$$\begin{aligned}\frac{\partial u}{\partial t}(x, t) &= \frac{\partial^2 u}{\partial x^2}(x, t) + F(u), \\ u(x, 0) &= f(x), \\ u(0, t) &= g(t), \\ u(1, t) &= h(t).\end{aligned}$$

Dividimos el intervalo $[0, 1]$ de x en n subintervalos iguales $[x_i, x_{i+1}]$, $i = 1, 2, \dots, n, n+1$ con,

$$x_i = x_{i-1} + \Delta x, x_1 = 0,$$

y $\Delta x = 1/n$ este es llamado el tamaño de la red y hacemos un proceso análogo para el intervalo de $[0, 1]$ de t no necesariamente del mismo tamaño de x . Sea Δt el tamaño de paso (en la dirección del tiempo). Los datos iniciales ahora son:

$$u_1^0 = f(x_1), u_2^0 = f(x_2), \dots, u_{n-1}^0 = f(x_{n-1}), u_n^0 = f(x_n),$$

y sea

$$u_i^j = u(x_i, t_j), t_j = t_{j-1} + \Delta t, x_i = x_{i-1} + \Delta x, j \geq 1, 1 \leq i \leq n+1.$$

Como la condición de frontera es conocida, entonces u_1^j y u_{n+1}^j son datos que ya se conocen para todo j y sólo debemos calcular para u_i^j para $i = 2, \dots, n$. Tomemos

$$v^j = [u_1^j, u_2^j, \dots, u_{n+1}^j].$$

La idea es dar el vector v^{j+1} a partir de v^j . Las ecuaciones diferenciales parciales pueden ser aproximadas por ecuaciones de diferencias teniendo las siguientes relaciones:

$$\begin{aligned}\frac{\partial u(x_i, t_j)}{\partial t} &\approx \frac{u_i^{j+1} - u_i^j}{\Delta t}, & \frac{\partial u(x_i, t_j)}{\partial x} &\approx \frac{u_{i+1}^j - u_i^j}{\Delta x}, \\ \frac{\partial^2 u(x_i, t_j)}{\partial x^2} &\approx \frac{u_{i+1}^j - 2u_i^j + u_{i-1}^j}{(\Delta x)^2}.\end{aligned}$$

Entonces la ecuación diferencial $u_t = u_{xx} + F(u)$ es ahora

$$\frac{u_i^{j+1} - u_i^j}{\Delta t} = \frac{u_{i+1}^j - 2u_i^j + u_{i-1}^j}{(\Delta x)^2} + F(u_i^j). \quad (2.11)$$

Sea

$$r = \frac{\Delta t}{(\Delta x)^2}, \quad c = 1 - 2r,$$

2. Ecuación del calor

donde $r < 1$ para que el método sea estable, ver [19], entonces la expresión (2.11) puede ser reescrita de la siguiente forma

$$u_i^{j+1} = ru_{i+1}^j + cu_i^j + ru_{i-1}^j + \Delta t F(u_i^j).$$

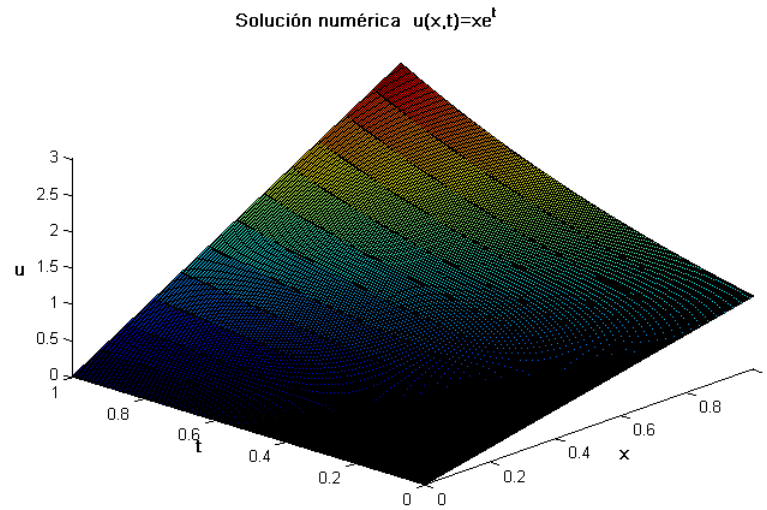
2.4.1. Ejemplos

Para la implementación en el computador utilizamos Matlab 7.8.0 (R2009a).

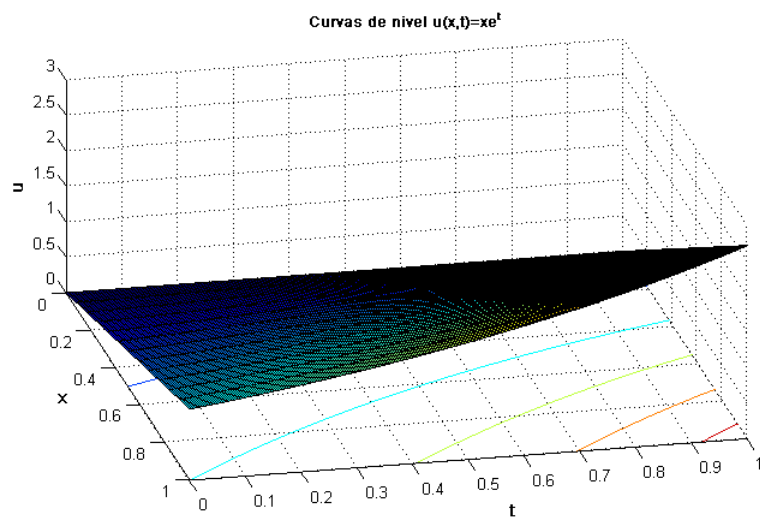
Ejemplo 1

$$\begin{aligned}\frac{\partial u}{\partial t}(x, t) &= \frac{\partial^2 u}{\partial x^2}(x, t) + F(u), \\ u(x, 0) &= x = f(x), \\ u(0, t) &= 0, \\ u(1, t) &= e^t = h(t),\end{aligned}$$

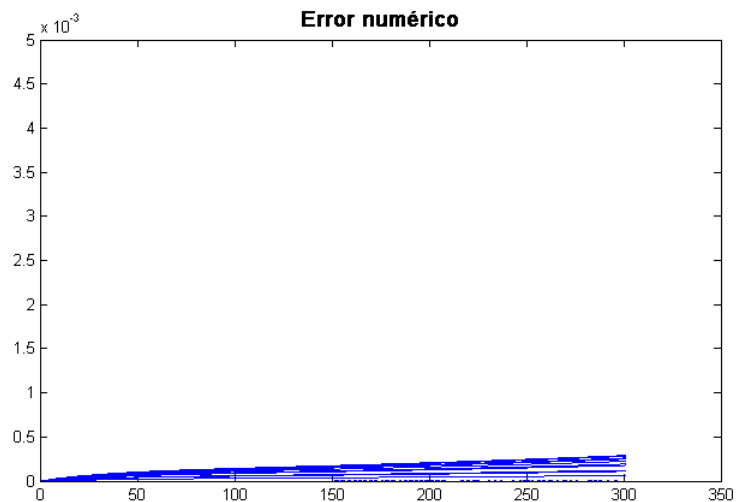
donde $F(u) = u$ y $u(x, t) = xe^t$.



2. Ecuación del calor



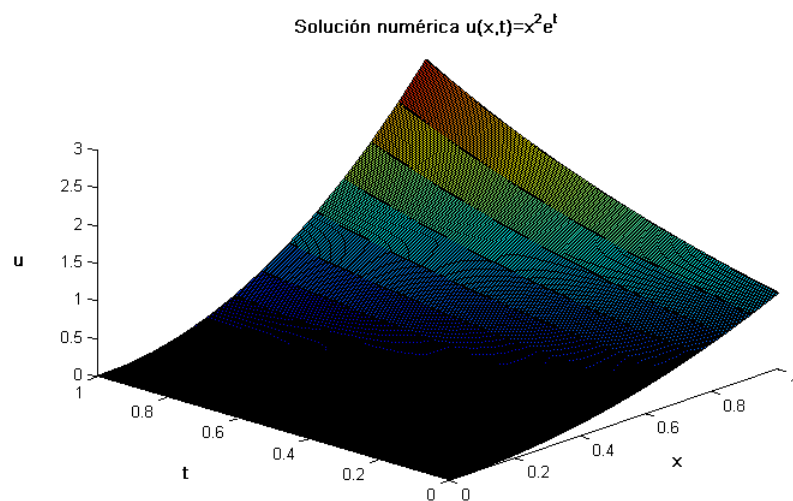
2. Ecuación del calor



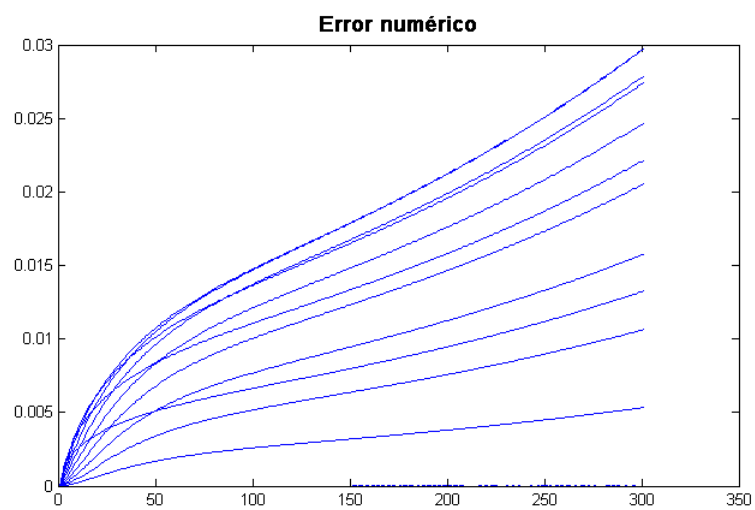
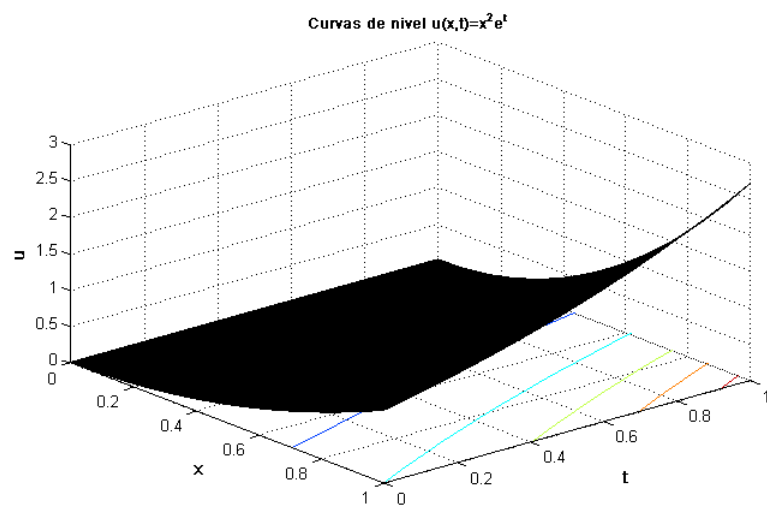
Ejemplo 2

$$\begin{aligned}\frac{\partial u}{\partial t}(x, t) &= \frac{\partial^2 u}{\partial x^2}(x, t) + F(x, t, u), \\ u(x, 0) &= x^2 = f(x), \\ u(0, t) &= 0, \\ u(1, t) &= e^t = h(t),\end{aligned}$$

donde $F(x, t, u) = u(x, t) - 2e(t)$ y $u(x, t) = x^2 e^t$.



2. Ecuación del calor



Capítulo 3

Regularización Tikhonov Generalizada

Sean X, Y, Z espacios de Hilbert adecuados con producto interno $\langle \cdot, \cdot \rangle_X, \langle \cdot, \cdot \rangle_Y, \langle \cdot, \cdot \rangle_Z$ respectivamente. Sean además $C \subseteq X$ un subconjunto acotado y los operadores $G : Y \longrightarrow Z$ y $\varphi : C \subset X \longrightarrow Y$ (lineales o no). Consideremos el siguiente problema implícito no lineal mal formulado

$$G \circ \varphi(x) = y \quad (3.1)$$

Suponemos que G y φ son operadores secuencialmente débilmente cerrados, es decir, para cualquier sucesión $\{x_n\} \subset C$ que converge débilmente a x^* entonces $\{\varphi(x_n)\}$ converge débilmente a φ^* y la sucesión $\{G(\varphi(x_n))\}$ converge débilmente a G^* en donde $\varphi(x^*) = \varphi^*$ y $G(\varphi(x^*)) = G^*$.

Para la noción de una solución del problema (3.1), elejimos el concepto de x^* - solución de norma mínima (x^* -m.n.s), x_0 la cual satisface

$$G \circ \varphi(x_0) = y,$$

y

$$\|x_0 - x^*\| = \min_x \{\|x - x^*\| : G \circ \varphi(x) = y\}.$$

En el problema (3.1) generalmente las soluciones no dependen continuamente de los datos (inestabilidad), esto es, pequeños cambios en los datos produce grandes cambios en la soluciones computadas. Generalmente, se recurre a métodos de regularización como el que emplearemos en este trabajo. Para investigar la solución del problema implícito introducimos el funcional

$$M_\alpha^\epsilon [G, \varphi, \varphi^*, x, x^*, y] = \|(G \circ \varphi)(x) - y\|_Z^2 + \epsilon \|\varphi(x) - \varphi^*\|_Y^2 + \alpha \|x - x^*\|_X^2, \quad (3.2)$$

3. Regularización Tikhonov Generalizada

donde ϵ, α son parámetros positivos, llamados parámetros de regularización. Y consideremos el problema de minimización

$$\min_{x \in D(\varphi)} M_\alpha^\epsilon[G, \varphi, \varphi^*, x, x^*, y]. \quad (3.3)$$

Este método se conoce como el método generalizado de Tikhonov que depende de dos parámetros de regularización para encontrar soluciones aproximadas a la solución deseada. A continuación presentaremos los teoremas de existencia, convergencia y estabilidad, de una manera general para espacios de Hilbert X, Y y Z , las pruebas que no se presentan aquí se pueden ver en [7], [15].

3.1. Existencia y monotonicidad

Teorema 3.1.1. *(Existencia.) Para todo $\alpha \neq 0$ y $\epsilon \neq 0$ existe una solución x_α^ϵ del funcional (3.3). Si $\varphi(C)$ y $G(\varphi(C))$ son acotados en Y y Z , respectivamente, entonces existe una solución $\hat{x}$ del problema (3.3) con $\alpha = \epsilon = 0$ la cual es una solución de mínimos cuadrados x^* -norma-mínima (x^* -s.m.c.m.n.), es decir,*

$$\|\hat{x} - x^*\|_X = \min_{x \in \Gamma} \{\|x - x^*\|_X : \Gamma = \min_x \|(G \circ \varphi)(x) - y\|_Z\}.$$

Demostración.

Consideremos el conjunto

$$S = \{(x, \varphi(x), (G \circ \varphi)(x)) : x \in C\} \subseteq X \times Y \times Z,$$

y definamos en este espacio el producto interno

$$\langle (x_1, y_1, z_1), (x_2, y_2, z_2) \rangle_{\epsilon, \alpha} = \alpha \langle x_1, x_2 \rangle_X + \epsilon \langle y_1, y_2 \rangle_Y + \langle z_1, z_2 \rangle_Z,$$

y la norma

$$\|(x_1, y_1, z_1)\|_{\epsilon, \alpha} = (\alpha \|x_1\|_X^2 + \epsilon \|y_1\|_Y^2 + \|z_1\|_Z^2)^{1/2}.$$

Para probar que S es un conjunto débilmente cerrado de $X \times Y \times Z$, sea $\{x_n\} \subseteq C$ tal que

$$(x_n, \varphi(x_n), (G \circ \varphi)(x_n)) \rightarrow (x_o, y_o, g_o),$$

esto implica que

$$\begin{aligned} x_n &\rightharpoonup x_o, \\ \varphi(x_n) &\rightharpoonup \varphi_o, \\ (G \circ \varphi)(x_n) &\rightharpoonup g_o. \end{aligned}$$

Como φ y G son operadores débilmente cerrados, $\varphi(x_o) = \varphi_o$ y $(G \circ \varphi)(x_o) = g_o$. Por lo tanto $(x_o, y_o, g_o) \in S$ y existe $(x_\alpha^\epsilon, \varphi(x_\alpha^\epsilon), (G \circ \varphi)(x_\alpha^\epsilon)) \in S$ tal que

$$\text{dis}(S, (x^*, \varphi^*, y)) = \|(x_\alpha^\epsilon - x^*, \varphi(x_\alpha^\epsilon) - \varphi^*, (G \circ \varphi)(x_\alpha^\epsilon) - y)\|_{\epsilon, \alpha},$$

3. Regularización Tikhonov Generalizada

o equivalentemente

$$M_\alpha^\epsilon[x_\alpha^\epsilon, y] = \min_{x \in C} M_\alpha^\epsilon[x, y].$$

Suponga ahora que $\varphi(C)$ y $G(\varphi(C))$ son acotados y considere que $\{\epsilon_n\}$ y $\{\alpha_n\}$ tales que $\epsilon_n \rightarrow 0$ y $\alpha_n \rightarrow 0$. Entonces existe una sucesión $\{x_n\} \subseteq C$ tal que

$$M_{\alpha_n}^{\epsilon_n}[x_n, y] = \min_{x \in C} M_{\alpha_n}^{\epsilon_n}[x, y],$$

o

$$\begin{aligned} \|(G \circ \varphi)(x_n) - y\|_Z^2 + \alpha_n \|x_n - x^*\|_X^2 + \epsilon_n \|\varphi(x_n) - \varphi^*\|_Y^2 \leq \\ \|(G \circ \varphi)(x) - y\|_Z^2 + \alpha_n \|x - x^*\|_X^2 + \epsilon_n \|\varphi(x) - \varphi^*\|_Y^2, \end{aligned}$$

para todo $x \in C$.

Además,

$$\overline{\lim_{n \rightarrow \infty}} \|(G \circ \varphi)(x_n) - y\|_Z \leq \|(G \circ \varphi)(x) - y\|_Z.$$

Dadas las sucesiones $\{x_n\}$, $\{\varphi(x_n)\}$ y $\{(G \circ \varphi)(x_n)\}$ acotadas, existe una subsucesión $\{x_{n_k}\}$ tal que

$$\begin{aligned} x_{n_k} &\rightharpoonup \hat{x}, \\ \varphi(x_{n_k}) &\rightharpoonup \varphi(\hat{x}), \\ (G \circ \varphi)(x_{n_k}) &\rightharpoonup (G \circ \varphi)(\hat{x}), \end{aligned}$$

entonces

$$\|(G \circ \varphi)(\hat{x}) - y\|_Z \leq \|(G \circ \varphi)(x) - y\|_Z,$$

para todo $x \in C$.

Esto significa que $\hat{x}$ es una solución del problema (3.3) con $\alpha = \epsilon = 0$. Además, $\hat{x}$ es también una solución de mínimos cuadrados de $(G \circ \varphi)(x) = y$. El conjunto de todas las soluciones de mínimos cuadrados es débilmente cerrado. De hecho, si denotamos el conjunto por L y asumimos $\{\hat{x}_n\} \subseteq L$, entonces para cualquier $x \in C$,

$$\|(G \circ \varphi)(\hat{x}_n) - y\|_Z \leq \|(G \circ \varphi)(x) - y\|_Z.$$

Además, existe una subsucesión $\{\hat{x}_{n_k}\}$ tal que

$$\begin{aligned} \hat{x}_{n_k} &\rightharpoonup w, \\ \varphi(\hat{x}_{n_k}) &\rightharpoonup \varphi(w), \\ (G \circ \varphi)(\hat{x}_{n_k}) &\rightharpoonup (G \circ \varphi)(w). \end{aligned}$$

Puesto que $\hat{x}_n \rightharpoonup w$, tenemos $w = w^*$, $(G \circ \varphi)(w) = (G \circ \varphi)(w^*)$ y

$$\|(G \circ \varphi)(w) - y\|_Z \leq \|(G \circ \varphi)(x) - y\|_Z,$$

3. Regularización Tikhonov Generalizada

para todo $x \in C$. Se deduce que $w \in L$, y además, L es un subconjunto débilmente cerrado de X y existe $x_0 \in L$ tal que

$$\|x_0 - x^*\|_X = \text{dis}(L, x^*).$$

Teorema 3.1.2. (*Monotonicidad*). Si x_α^ϵ es una solución del funcional (3.3), entonces

i) Si $\epsilon_2 > \epsilon_1 \geq 0$, entonces

$$\sup \|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y \leq \inf \|\varphi(x_\alpha^{\epsilon_1}) - \psi^*\|_Y$$

y

$$\sup \{ \|(G \circ \varphi)(x_\alpha^{\epsilon_1}) - y\|_Z^2 + \alpha \|x_\alpha^{\epsilon_1} - x^*\|_X^2 \} \leq \inf \{ \|(G \circ \varphi)(x_\alpha^{\epsilon_2}) - y\|_Z^2 + \alpha \|x_\alpha^{\epsilon_2} - x^*\|_X^2 \},$$

para todo $\alpha \geq 0$.

ii) Si el problema no regularizado tiene una solución $\hat{x}$, entonces

$$\begin{aligned} \|(G \circ \varphi)(x_\alpha^{\epsilon_2}) - y\|_Z^2 + \alpha \|x_\alpha^{\epsilon_2} - x^*\|_X^2 + \epsilon_2 \|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y^2 \leq \\ \|(G \circ \varphi)(\hat{x}) - y\|_Z^2 + \alpha \|\hat{x} - x^*\|_X^2 + \epsilon_2 \|\varphi(\hat{x}) - \psi^*\|_Y^2, \end{aligned}$$

para todo $\alpha \geq 0$.

iii) Si $\alpha_2 > \alpha_1 \geq 0$, entonces

$$\sup \|x_{\alpha_2}^\epsilon - x^*\|_X \leq \inf \|x_{\alpha_1}^\epsilon - x^*\|_X,$$

y

$$\begin{aligned} \sup \{ \|(G \circ \varphi)(x_{\alpha_1}^\epsilon) - y\|_Z^2 + \epsilon \|\varphi(x_{\alpha_1}^\epsilon) - \psi^*\|_Y^2 \} \\ \leq \inf \{ \|(G \circ \varphi)(x_{\alpha_2}^\epsilon) - y\|_Z^2 + \epsilon \|\varphi(x_{\alpha_2}^\epsilon) - \psi^*\|_Y^2 \}, \end{aligned}$$

para todo $\epsilon \geq 0$.

iv) Si $\alpha_2 = \epsilon_2$, $\alpha_1 = \epsilon_1$ y $\epsilon_2 > \epsilon_1 \geq 0$, entonces

$$\sup \{ \|x^{\epsilon_2} - x^*\|_X^2 + \|\varphi(x^{\epsilon_2}) - \psi^*\|_Y^2 \}^{1/2} \leq \inf \{ \|x^{\epsilon_1} - x^*\|_X^2 + \|\varphi(x^{\epsilon_1}) - \psi^*\|_Y^2 \}^{1/2}$$

y

$$\sup \|(G \circ \varphi)(x^{\epsilon_1}) - y\|_Z \leq \inf \|(G \circ \varphi)(x^{\epsilon_2}) - y\|_Z.$$

3. Regularización Tikhonov Generalizada

Demostración:

i) Sumando $(\epsilon_2 - \epsilon_1) \|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y^2$ a ambos lados de la desigualdad

$$\begin{aligned} \|(G \circ \varphi)(x_\alpha^{\epsilon_1}) - y\|_Z^2 + \alpha \|x_\alpha^{\epsilon_1} - x^*\|_X^2 + \epsilon_1 \|\varphi(x_\alpha^{\epsilon_1}) - \psi^*\|_Y^2 \leq \\ \|(G \circ \varphi)(x_\alpha^{\epsilon_2}) - y\|_Z^2 + \alpha \|x_\alpha^{\epsilon_2} - x^*\|_X^2 + \epsilon_1 \|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y^2, \end{aligned} \quad (3.4)$$

tenemos

$$\begin{aligned} \|(G \circ \varphi)(x_\alpha^{\epsilon_1}) - y\|_Z^2 + \alpha \|x_\alpha^{\epsilon_1} - x^*\|_X^2 + \epsilon_1 \|\varphi(x_\alpha^{\epsilon_1}) - \psi^*\|_Y^2 + (\epsilon_2 - \epsilon_1) \|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y^2 \\ \leq \|(G \circ \varphi)(x_\alpha^{\epsilon_2}) - y\|_Z^2 + \alpha \|x_\alpha^{\epsilon_2} - x^*\|_X^2 + \epsilon_1 \|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y^2 + (\epsilon_2 - \epsilon_1) \|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y^2, \end{aligned}$$

cancelando términos semejantes obtenemos

$$\begin{aligned} \|(G \circ \varphi)(x_\alpha^{\epsilon_1}) - y\|_Z^2 + \alpha \|x_\alpha^{\epsilon_1} - x^*\|_X^2 + \epsilon_2 \|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y^2 \\ + \epsilon_1 (\|\varphi(x_\alpha^{\epsilon_1}) - \psi^*\|_Y - \|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y) \\ \leq \|(G \circ \varphi)(x_\alpha^{\epsilon_2}) - y\|_Z^2 + \alpha \|x_\alpha^{\epsilon_2} - x^*\|_X^2 + \epsilon_2 \|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y^2 \\ \leq \|(G \circ \varphi)(x_\alpha^{\epsilon_1}) - y\|_Z^2 + \alpha \|x_\alpha^{\epsilon_1} - x^*\|_X^2 + \epsilon_2 \|\varphi(x_\alpha^{\epsilon_1}) - \psi^*\|_Y^2, \end{aligned}$$

luego

$$\epsilon_2 (\|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y^2 - \|\varphi(x_\alpha^{\epsilon_1}) - \psi^*\|_Y^2) + \epsilon_1 (\|\varphi(x_\alpha^{\epsilon_1}) - \psi^*\|_Y^2 - \|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y^2) \leq 0$$

y

$$(\epsilon_2 - \epsilon_1) (\|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y^2 - \|\varphi(x_\alpha^{\epsilon_1}) - \psi^*\|_Y^2) \leq 0.$$

Puesto que $\epsilon_2 > \epsilon_1$, tenemos

$$\|\varphi(x_\alpha^{\epsilon_2}) - \psi^*\|_Y^2 \leq \|\varphi(x_\alpha^{\epsilon_1}) - \psi^*\|_Y^2.$$

Esta última inecuación y (3.4) implican

$$\|(G \circ \varphi)(x_\alpha^{\epsilon_1}) - y\|_Z^2 + \alpha \|x_\alpha^{\epsilon_1} - x^*\|_X^2 \leq \|(G \circ \varphi)(x_\alpha^{\epsilon_2}) - y\|_Z^2 + \alpha \|x_\alpha^{\epsilon_2} - x^*\|_X^2,$$

lo cual demuestra la segunda parte de i).

- ii) Es una consecuencia directa del hecho que $(x_\alpha^{\epsilon_2})$ es una solución del problema (3.3).
- iii) Usamos la misma idea utilizada en la demostración de la parte i). Dado que $(x_{\alpha_i}^\epsilon)$ es una solución del problema (3.3), podemos escribir

$$\begin{aligned} \|(G \circ \varphi)(x_{\alpha_1}^\epsilon) - y\|_Z^2 + \alpha_1 \|x_{\alpha_1}^\epsilon - x^*\|_X^2 + \epsilon \|\varphi(x_{\alpha_1}^\epsilon) - \psi^*\|_Y^2 \leq \\ \|(G \circ \varphi)(x_{\alpha_2}^\epsilon) - y\|_Z^2 + \alpha_1 \|x_{\alpha_2}^\epsilon - x^*\|_X^2 + \epsilon \|\varphi(x_{\alpha_2}^\epsilon) - \psi^*\|_Y^2. \end{aligned} \quad (3.5)$$

3. Regularización Tikhonov Generalizada

Sumando $(\alpha_2 - \alpha_1) \|x_{\alpha_2}^\epsilon - x^*\|_X$ a ambos lados de la inecuación anterior, obtenemos

$$\begin{aligned}
\|(G \circ \varphi)(x_{\alpha_1}^\epsilon) - y\|_Z^2 &+ \alpha_2 \|x_{\alpha_2}^\epsilon - x^*\|_X^2 + \alpha_1 (\|x_{\alpha_1}^\epsilon - x^*\|_X^2 - \|x_{\alpha_2}^\epsilon - x^*\|_X^2) \\
&+ \epsilon \|\varphi(x_{\alpha_1}^\epsilon) - \psi^*\|_Y^2 \\
&\leq \|(G \circ \varphi)(x_{\alpha_2}^\epsilon) - y\|_Z^2 + \alpha_2 \|x_{\alpha_2}^\epsilon - x^*\|_X^2 + \epsilon \|\varphi(x_{\alpha_2}^\epsilon) - \psi^*\|_Y^2 \\
&\leq \|(G \circ \varphi)(x_{\alpha_1}^\epsilon) - y\|_Z^2 + \alpha_2 \|x_{\alpha_1}^\epsilon - x^*\|_X^2 + \epsilon \|\varphi(x_{\alpha_1}^\epsilon) - \psi^*\|_Y^2,
\end{aligned}$$

y agrupando términos tenemos

$$\alpha_2 (\|x_{\alpha_2}^\epsilon - x^*\|_X^2 - \|x_{\alpha_1}^\epsilon - x^*\|_X^2) + \alpha_1 (\|x_{\alpha_1}^\epsilon - x^*\|_X^2 - \|x_{\alpha_2}^\epsilon - x^*\|_X^2) \leq 0.$$

De este modo

$$(\alpha_2 - \alpha_1) (\|x_{\alpha_2}^\epsilon - x^*\|_X^2 - \|x_{\alpha_1}^\epsilon - x^*\|_X^2) \leq 0,$$

como $\alpha_2 > \alpha_1$ se tiene que

$$\|x_{\alpha_2}^\epsilon - x^*\|_X^2 \leq \|x_{\alpha_1}^\epsilon - x^*\|_X^2.$$

La última desigualdad y (3.5) implican la segunda parte de iii). Es decir,

$$\begin{aligned}
\|(G \circ \varphi)(x_{\alpha_1}^\epsilon) - y\|_Z^2 &+ \alpha_1 \|x_{\alpha_1}^\epsilon - x^*\|_X^2 + \epsilon \|\varphi(x_{\alpha_1}^\epsilon) - \psi^*\|_Y^2 \\
&\leq \|(G \circ \varphi)(x_{\alpha_2}^\epsilon) - y\|_Z^2 + \alpha_1 \|x_{\alpha_2}^\epsilon - x^*\|_X^2 + \epsilon \|\varphi(x_{\alpha_2}^\epsilon) - \psi^*\|_Y^2 \\
&\leq \|(G \circ \varphi)(x_{\alpha_2}^\epsilon) - y\|_Z^2 + \alpha_1 \|x_{\alpha_1}^\epsilon - x^*\|_X^2 + \epsilon \|\varphi(x_{\alpha_2}^\epsilon) - \psi^*\|_Y^2,
\end{aligned}$$

de donde

$$\|(G \circ \varphi)(x_{\alpha_1}^\epsilon) - y\|_Z^2 + \epsilon \|\varphi(x_{\alpha_1}^\epsilon) - \psi^*\|_Y^2 \leq \|(G \circ \varphi)(x_{\alpha_2}^\epsilon) - y\|_Z^2 + \epsilon \|\varphi(x_{\alpha_2}^\epsilon) - \psi^*\|_Y^2.$$

iv) Para esta prueba, sumamos $(\epsilon_2 - \epsilon_1) (\|x^{\epsilon_2} - x^*\|_X^2 + \|\varphi(x^{\epsilon_2}) - \psi^*\|_Y^2)$ a la desigualdad

$$\begin{aligned}
\|(G \circ \varphi)(x^{\epsilon_1}) - y\|_Z^2 + \epsilon_1 \|x^{\epsilon_1} - x^*\|_X^2 + \epsilon_1 \|\varphi(x^{\epsilon_1}) - \psi^*\|_Y^2 &\leq \\
\|(G \circ \varphi)(x^{\epsilon_2}) - y\|_Z^2 + \epsilon_1 \|x^{\epsilon_2} - x^*\|_X^2 + \epsilon_1 \|\varphi(x^{\epsilon_2}) - \psi^*\|_Y^2, &\quad (3.6)
\end{aligned}$$

y se obtiene

$$\begin{aligned}
\|(G \circ \varphi)(x^{\epsilon_1}) - y\|_Z^2 &+ \epsilon_2 (\|x^{\epsilon_2} - x^*\|_X^2 + \|\varphi(x^{\epsilon_2}) - \psi^*\|_Y^2) \\
&+ \epsilon_1 (\|x^{\epsilon_1} - x^*\|_X^2 - \|x^{\epsilon_2} - x^*\|_X^2 + \|\varphi(x^{\epsilon_1}) - \psi^*\|_Y^2 \\
&- \|\varphi(x^{\epsilon_2}) - \psi^*\|_Y^2) \\
&\leq \|(G \circ \varphi)(x^{\epsilon_2}) - y\|_Z^2 + \epsilon_2 \|x^{\epsilon_2} - x^*\|_X^2 + \epsilon_2 \|\varphi(x^{\epsilon_2}) - \psi^*\|_Y^2 \\
&\leq \|(G \circ \varphi)(x^{\epsilon_1}) - y\|_Z^2 + \epsilon_2 \|x^{\epsilon_1} - x^*\|_X^2 + \epsilon_2 \|\varphi(x^{\epsilon_1}) - \psi^*\|_Y^2.
\end{aligned}$$

3. Regularización Tikhonov Generalizada

Por lo tanto,

$$\begin{aligned} & \epsilon_2 (\|x^{\epsilon_2} - x^*\|_X^2 - \|x^{\epsilon_1} - x^*\|_X^2 + \|\varphi(x^{\epsilon_2}) - \psi^*\|_Y^2 - \|\varphi(x^{\epsilon_1}) - \psi^*\|_Y^2) \\ & - \epsilon_1 (\|x^{\epsilon_2} - x^*\|_X^2 - \|x^{\epsilon_1} - x^*\|_X^2 + \|\varphi(x^{\epsilon_2}) - \psi^*\|_Y^2 - \|\varphi(x^{\epsilon_1}) - \psi^*\|_Y^2) \leq 0 \end{aligned}$$

y

$$(\epsilon_2 - \epsilon_1) (\|x^{\epsilon_2} - x^*\|_X^2 - \|x^{\epsilon_1} - x^*\|_X^2 + \|\varphi(x^{\epsilon_2}) - \psi^*\|_Y^2 - \|\varphi(x^{\epsilon_1}) - \psi^*\|_Y^2) \leq 0.$$

Como $\epsilon_2 > \epsilon_1$, tenemos que

$$\|x^{\epsilon_2} - x^*\|_X^2 + \|\varphi(x^{\epsilon_2}) - \psi^*\|_Y^2 \leq \|x^{\epsilon_1} - x^*\|_X^2 + \|\varphi(x^{\epsilon_1}) - \psi^*\|_Y^2.$$

De la última desigualdad y (3.6) tenemos

$$\begin{aligned} \|(G \circ \varphi)(x^{\epsilon_1}) - y\|_Z^2 &+ \epsilon_1 \|x^{\epsilon_1} - x^*\|_X^2 + \epsilon_1 \|\varphi(x^{\epsilon_1}) - \psi^*\|_Y^2 \\ &\leq \|(G \circ \varphi)(x^{\epsilon_2}) - y\|_Z^2 + \epsilon_1 \|x^{\epsilon_2} - x^*\|_X^2 + \epsilon_1 \|\varphi(x^{\epsilon_2}) - \psi^*\|_Y^2 \\ &\leq \|(G \circ \varphi)(x^{\epsilon_2}) - y\|_Z^2 + \epsilon_1 \|x^{\epsilon_1} - x^*\|_X^2 + \epsilon_1 \|\varphi(x^{\epsilon_1}) - \psi^*\|_Y^2, \end{aligned}$$

cancelando términos comunes tenemos

$$\|(G \circ \varphi)(x^{\epsilon_1}) - y\|_Z^2 \leq \|(G \circ \varphi)(x^{\epsilon_2}) - y\|_Z^2.$$

3.2. Estabilidad y convergencia

Teorema 3.2.1. (*Estabilidad*). Sean $\epsilon > 0$ y $\alpha > 0$ fijos. Sea $\{y_n\}$ una sucesión en Y tal que $y_n \rightarrow y$ en Y . Sea $\{x_n\}$ una sucesión de soluciones del problema (3.3) para el funcional $M_\alpha^\epsilon[x, y_n]$. Si $\varphi(C)$ y $G(\varphi(C))$ son acotados en Y y Z , respectivamente, entonces existe una subsucesión, $\{x_{n_k}\}$, que converge débilmente, tal que cada límite débil w de la subsucesión es una solución de (3.3), además $\varphi(x_{n_k}) \rightarrow \varphi(w)$ y $(G \circ \varphi)(x_{n_k}) \rightarrow (G \circ \varphi)(w)$.

Demostración. Se puede ver en [7].

Lema 3.2.2. Suponga que $\varphi(C)$ y $G(\varphi(C))$ son acotados en Y y Z , respectivamente. Sea $\alpha_j \rightarrow 0$ y $\epsilon_n \rightarrow 0$. Entonces existe una subsucesión convergente en C cuyo límite es una solución del problema (P_α^ϵ) . Para $\alpha = \epsilon = 0$.

Demostración.

Sea $\{x_j^n\}$ una solución (3.3). Puesto que $\{x_j^n\}$ y $\{\varphi(x_j^n)\}$ son sucesiones acotadas en C y $\varphi(C)$, respectivamente, para j fijo tenemos $x_j^{n_k} \rightharpoonup \omega_j$, $\varphi(x_j^{n_k}) \rightharpoonup \varphi(\omega_j)$, y $(G \circ \varphi)(x_j^{n_k}) \rightharpoonup (G \circ \varphi)(\omega_j)$. Es claro que $\{\omega_j\}$ es acotada y

$$\begin{aligned} \|(G \circ \varphi)(x_j^{n_k}) - y\|_Z^2 + \alpha_j \|x_j^{n_k} - x^*\|_X^2 + \epsilon_{n_k} \|\varphi(x_j^{n_k}) - \psi^*\|_Y^2 &\leq \\ \|(G \circ \varphi)(x) - y\|_Z^2 + \alpha_j \|x - x^*\|_X^2 + \epsilon_{n_k} \|\varphi(x) - \psi^*\|_Y^2, \end{aligned}$$

3. Regularización Tikhonov Generalizada

para todo $x \in C$.

Sacando lím ínf tenemos

$$\|(G \circ \varphi)(\omega_j) - y\|_Z^2 + \alpha_j \|\omega_j - x^*\|_X^2 \leq \|(G \circ \varphi)(x) - y\|_Z^2 + \alpha_j \|x - x^*\|_X^2,$$

para todo $x \in C$.

Así, ω_j es una solución del problema (P_α^ϵ) , para $\alpha = \alpha_j, \epsilon = 0$. Puesto que $\{\omega_j\}$ es acotada, existe una subsucesión $\{\omega_{j_k}\}$ tal que

$$\omega_{j_k} \rightharpoonup \omega, \quad \varphi(\omega_{j_k}) \rightharpoonup \varphi(\omega) \text{ y } (G \circ \varphi)(\omega_{j_k}) \rightharpoonup (G \circ \varphi)(\omega).$$

También

$$\|(G \circ \varphi)(\omega_{j_k}) - y\|_Z^2 + \alpha_{j_k} \|\omega_{j_k} - x^*\|_X^2 \leq \|(G \circ \varphi)(x) - y\|_Z^2 + \alpha_{j_k} \|x - x^*\|_X^2,$$

para todo $x \in C$.

Sacando lím ínf, tenemos

$$\|(G \circ \varphi)(\omega) - y\|_Z^2 \leq \|(G \circ \varphi)(x) - y\|_Z^2,$$

para todo $x \in C$.

En consecuencia, ω es una solución del problema (3.3) $\alpha = \epsilon = 0$. Finalmente, del teorema 3.1.2 iii), tomando $\alpha_1 = 0$ y $\epsilon = 0$,

$$\overline{\lim} \|\omega_{j_k} - x^*\|_X \leq \|\omega - x^*\|_X,$$

y

$$\omega_{j_k} \rightarrow \omega.$$

Teorema 3.2.3. (Convergencia-estabilidad). Sea $\varphi(C)$ y $F(\varphi(x))$ acotados en Y y Z , respectivamente.

i) Suponga que el problema $(G \circ \varphi)(x) = y$ tiene una solución y

$$\|y - y^{\delta_n}\|_Z \leq \delta_n \rightarrow 0, \quad \alpha_n \rightarrow 0, \quad \epsilon_n \rightarrow 0, \quad \frac{\delta_n^2}{\alpha_n} \rightarrow 0 \text{ y } \frac{\alpha_n}{\epsilon_n} \rightarrow 0.$$

Entonces la sucesión $\{x_{\alpha_n}^{\epsilon_n}\}$ tiene una subsucesión convergente fuertemente a x^* -s.m.c.m.n., $\hat{x}$, tal que $(G \circ \varphi)(\hat{x}) = y$ y $(G \circ \varphi)(x_{n_k}) \rightarrow (G \circ \varphi)(\hat{x}) = y$.

3. Regularización Tikhonov Generalizada

ii) Suponga que el problema $(G \circ \varphi)(x) = y$ tiene una solución y

$$\|y - y^{\delta_n}\|_Z \leq \delta_n \rightarrow 0, \alpha_n \rightarrow 0, \epsilon_n \rightarrow 0, \frac{\delta_n^2}{\epsilon_n} \rightarrow 0 \text{ y } \frac{\alpha_n}{\epsilon_n} \rightarrow 0.$$

Entonces existe una subsucesión de $\{x_{\alpha_n}^{\epsilon_n}\}$, $\{x_{\alpha_{n_k}}^{\epsilon_{n_k}}\}$ que converge débilmente y cada límite débil $\bar{x}$ de $\{x_{\alpha_{n_k}}^{\epsilon_{n_k}}\}$ satisface $(G \circ \varphi)(\bar{x}) = y$, $\varphi(x_{\alpha_{n_k}}^{\epsilon_{n_k}}) \rightarrow \varphi(\bar{x})$ y $(G \circ \varphi)(x_{\alpha_{n_k}}^{\epsilon_{n_k}}) \rightarrow (G \circ \varphi)(\bar{x}) = y$.

iii) Suponga que el problema $(G \circ \varphi)(x) = y$ tiene una solución y

$$\|y - y^{\delta_n}\|_Z \leq \delta_n \rightarrow 0, \alpha_n = \epsilon_n \rightarrow 0, \text{ y } \frac{\delta_n^2}{\alpha_n} \rightarrow 0.$$

Entonces la sucesión $\{x_{\alpha_n}^{\epsilon_n}\}$ tiene una subsucesión, $\{x_{\alpha_{n_k}}^{\epsilon_{n_k}}\}$, que converge fuertemente a $\omega \in C$ tal que $(G \circ \varphi)(\omega) = y$, $\varphi(x_{\alpha_{n_k}}^{\epsilon_{n_k}}) \rightarrow \varphi(\omega)$ y $(G \circ \varphi)(x_{\alpha_{n_k}}^{\epsilon_{n_k}}) \rightarrow (G \circ \varphi)(\omega) = y$. Más aún, $(\omega, \varphi(\omega))$ es una (x^*, ψ^*) -s.m.n.

Demostración. Se puede ver en [7].

Lema 3.2.4. Dados a, b y c números reales positivos, si $a^2 \leq ab + c^2$ entonces $a \leq b + c$.

Demostración:

Supongamos que $a > b + c$, entonces $c < a - b$, como $a^2 \leq ab + c^2 < ab + (a - b)^2$ y $a^2 < ab + a^2 - 2ab + b^2$ luego $0 < b(b - a)$, y como $b > 0$ entonces $b - a > 0$ y de aquí $b > a$. y se tiene que $a^2 \leq ab + c^2 < b^2 + c^2$, luego $a^2 < (b + c)^2$, entonces $a < b + c$; lo que es contradictorio porque se ha supuesto que $a > b + c$. En consecuencia $a \leq b + c$. [15].

Teorema 3.2.5. (Razón de Convergencia). Sea C convexo, $\|y - y^\delta\|_Z \leq \delta$, y $(G \circ \varphi)(x_0) = y$. Además, se tienen las siguientes condiciones:

a. G y φ son diferenciables en el sentido de Fréchet.

b. $(G \circ \varphi)'(x_0)(x_\alpha^\delta - x_0) = (G \circ \varphi)(x_\alpha^\delta) - (G \circ \varphi)(x_0) - \hat{r}_\alpha = (G \circ \varphi)(x_\alpha^\delta) - y - \hat{r}_\alpha$, donde

$$\|\hat{r}_\alpha\|_Z \leq \frac{L_0}{2} \|x_\alpha^\delta - x_0\|_X^2.$$

c. $\varphi(x_\alpha^\delta) - \varphi(x_0) = \varphi'(x_0)(x_\alpha^\delta - x_0) + \hat{r}_\alpha$, donde $\|\hat{r}_\alpha\|_Z \leq \frac{L_1}{2} \|x_\alpha^\delta - x_0\|_X^2$.

d. Existe $\mu \in Z$ talque $(x^* - x_0) + [\varphi'(x_0)]^* v = [(G \circ \varphi)(x_0)]^* \mu$, donde $v = \psi^* - \varphi(x_0)$.

e. $L = L_0 \|\mu\|_Z + L_1 \|v\|_Y$ y $L < 1$. Entonces, con $\alpha = \epsilon = \delta$,

$$\|x_\alpha^\delta - x_0\|_X \leq \vartheta(\sqrt{\alpha}).$$

3. Regularización Tikhonov Generalizada

Demostración:

Usando $x = x_0$ y $\|y - y^\delta\|_Z \leq \delta$,

$$\begin{aligned} \|(G \circ \varphi)(x_\alpha) - y^\delta\|_Z^2 + \alpha \|x_\alpha - x^*\|_X^2 + \alpha \|\varphi(x_\alpha) - \psi^*\|_Y^2 \leq \\ \delta^2 + \alpha \|x_0 - x^*\|_X^2 + \alpha \|\varphi(x_\alpha) - \psi^*\|_Y^2. \end{aligned}$$

De

$$\|x_\alpha - x^* + x_0 - x_0\|_X^2 = \|(x_\alpha - x_0) - (x^* - x_0)\|_X^2,$$

tenemos

$$\|x_\alpha - x^*\|_X^2 = \|x_\alpha - x^*\|_X^2 + \|x_0 - x^*\|_X^2 - 2 \langle x_\alpha - x_0, x^* - x_0 \rangle,$$

y

$$\|\varphi(x_\alpha) - \psi^*\|_Y^2 = \|\varphi(x_\alpha) - \varphi(x_0)\|_Y^2 + \|\varphi(x_0) - \psi^*\|_Y^2 - 2 \langle \varphi(x_\alpha) - \varphi(x_0), \psi^* - \varphi(x_0) \rangle.$$

Entonces

$$\begin{aligned} & \|(G \circ \varphi)(x_\alpha) - y^\delta\|_Z^2 + \alpha \|x_\alpha - x_0\|_X^2 + \alpha \|\varphi(x_\alpha) - \varphi(x_0)\|_Y^2 \\ & \leq \delta^2 + 2\alpha \langle x_\alpha - x_0, x^* - x_0 \rangle + 2\alpha \langle \varphi(x_\alpha) - \varphi(x_0), \psi^* - \varphi(x_0) \rangle \\ & = \delta^2 + 2\alpha \langle x_\alpha - x_0, x^* - x_0 \rangle + 2\alpha \langle \varphi'(x_0)(x_\alpha - x_0) + r_\alpha, v \rangle \\ & = \delta^2 + 2\alpha \langle x_\alpha - x_0, x^* - x_0 \rangle + 2\alpha \langle (x_\alpha - x_0), \varphi'(x_0) * v \rangle + 2\alpha \langle r_\alpha, v \rangle \\ & = \delta^2 + 2\alpha \langle x_\alpha - x_0, x^* - x_0 + \varphi'(x_0) * v \rangle + 2\alpha \langle r_\alpha, v \rangle \\ & = \delta^2 + 2\alpha \langle (G \circ \varphi)'(x_0)(x_\alpha - x_0), \mu \rangle + 2\alpha \langle r_\alpha, v \rangle \\ & = \delta^2 + 2\alpha \langle (G \circ \varphi)(x_\alpha) - y, \mu \rangle - 2\alpha \langle \hat{r}_\alpha, \mu \rangle + 2\alpha \langle r_\alpha, v \rangle \\ & = \delta^2 + 2\alpha \langle (G \circ \varphi)(x_\alpha) - y^\delta, \mu \rangle + 2\alpha \langle y^\delta - y, \mu \rangle - 2\alpha \langle \hat{r}_\alpha, \mu \rangle + 2\alpha \langle r_\alpha, v \rangle \\ & \leq \delta^2 + 2\alpha \|\mu\|_Z \|(G \circ \varphi)(x_\alpha) - y^\delta\|_Z + 2\alpha \|\mu\|_Z \delta + (L_0 \|\mu\|_Z + L_1 \|v\|_Y) \alpha \|x_\alpha - x_0\|_Y^2 \\ & = \delta^2 + 2\alpha \|\mu\|_Z \|(G \circ \varphi)(x_\alpha) - y^\delta\|_Z + 2\alpha \|\mu\|_Z \delta + \alpha L \|x_\alpha - x_0\|_Y^2, \end{aligned}$$

luego,

$$\begin{aligned} \|(G \circ \varphi)(x_\alpha) - y^\delta\|_Z^2 + \alpha(1 - L) \|x_\alpha - x_0\|_X^2 + \alpha \|\varphi(x_\alpha) - \varphi(x_0)\|_Y^2 \leq \\ \delta^2 + 2\alpha\delta \|\mu\| + 2\alpha \|\mu\|_Z \|(G \circ \varphi)(x_\alpha) - y^\delta\|_Z, \quad (3.7) \end{aligned}$$

y de esto es claro que

$$\|(G \circ \varphi)(x_\alpha) - y^\delta\|_Z^2 \leq \delta^2 + 2\alpha\delta \|\mu\| + 2\alpha \|\mu\|_Z \|(G \circ \varphi)(x_\alpha) - y^\delta\|_Z.$$

Usando el lema anterior tenemos

$$\|(G \circ \varphi)(x_\alpha) - y^\delta\|_Z \leq (\delta^2 + 2\alpha\delta \|\mu\|_Z)^{1/2} + 2\alpha \|\mu\|_Z.$$

Ahora de (3.7) y de la última estimación tenemos que

3. Regularización Tikhonov Generalizada

$$\begin{aligned}\alpha(1-L) \|x_\alpha - x_0\|_X^2 &\leq \delta^2 + 2\alpha\delta \|\mu\| + 2\alpha \|\mu\|_Z ((\delta^2 + 2\alpha\delta \|\mu\|_Z)^{1/2} + 2\alpha \|\mu\|_Z) \\ &= \delta^2 + 2\alpha\delta \|\mu\| + 2\alpha \|\mu\|_Z (\delta^2 + 2\alpha\delta \|\mu\|_Z)^{1/2} + 4\alpha^2 \|\mu\|_Z \\ &\leq [(\delta^2 + 2\alpha\delta \|\mu\|_Z)^{1/2} + 2\alpha \|\mu\|_Z]^2\end{aligned}$$

y

$$\|x_\alpha - x_0\|_X \leq \frac{(\delta^2 + 2\alpha\delta \|\mu\|_Z)^{1/2} + 2\alpha \|\mu\|_Z}{\sqrt{\alpha}\sqrt{1-L}} = \vartheta(\sqrt{\alpha}).$$

Capítulo 4

Fórmulas numéricas

4.1. Formulación variacional y operadores débilmente cerrados

En la formulación variacional del problema:

$$u_t(x, t) = u_{xx}(x, t) + F(x, t, u), \quad (x, t) \in [0, 1] \times [0, 1],$$

bajo las condiciones

$$u(x, 0) = f(x) \quad 0 \leq x \leq 1,$$

$$u(1, t) = g(t) \quad 0 \leq t \leq 1,$$

$$u(0, t) = h(t) \quad 0 \leq t \leq 1.$$

es necesario determinar $u \in H^1(\Omega)$, donde $\Omega = [0, 1] \times [0, 1]$; tal que

$$\int_0^1 \frac{\partial u}{\partial t} \phi \, dx = \int_0^1 \frac{\partial u}{\partial x} \frac{\partial \phi}{\partial x} \, dx + \int_0^1 F(x, t, u) \phi \, dx, \quad \forall \phi \in H_0^1(\Omega) \quad (4.1)$$

y

$$u(x, 0) = f(x)$$

Estamos interesados en identificar el término fuente F cuando tenemos datos observados u dados por

$$u^\delta(x, t) = u(x, t) + \text{ruido}, \quad t \in (t_1, t_2),$$

tal que $\|u^\delta - u\|_{L_2(Q_T)} \leq \delta$, donde $Q_T = [0, 1] \times [t_1, t_2]$. Esta clase de problemas de identificación son usualmente *mal puestos* en el sentido en que pequeñas perturbaciones en los datos pueden producir grandes errores en las soluciones calculadas. Es decir, que en el intervalo (t_1, t_2) donde u se conoce con errores debido a la práctica, debemos encontrar u_t y u_{xx} ; diferenciación numérica que resulta ser un problema *inverso mal puesto*.

4. Fórmulas numéricas

Para estabilizar este problema, proponemos el método de regularización generalizado de Tikhonov definido en el capítulo anterior, que depende de dos parámetros de regularización α y ϵ donde G y φ son operadores definidos a continuación.

Sea

$$\begin{aligned}\varphi : C_1^2(\Omega) \subset H^1(\Omega) &\longrightarrow H^1(\Omega), \\ \varphi(F)(x, t) &= u(F, x, t),\end{aligned}\tag{4.2}$$

donde u es la solución del problema que depende de F , y

$$\begin{aligned}G : H^1(\Omega) &\longrightarrow H^1([0, 1]) \times L^2([0, 1]), \\ G(u)(x) &= \left(\int_{t_1}^{t_2} u(x, t) dt, \int_{t_1}^{t_2} \frac{\partial u(x, t)}{\partial x} dt \right).\end{aligned}\tag{4.3}$$

Por lo tanto definimos el operador implícito $G \circ \varphi$ como

$$G(\varphi(F)) = \left(\int_{t_1}^{t_2} u(F, x, t) dt, \int_{t_1}^{t_2} \frac{\partial u(F, x, t)}{\partial x} dt \right).\tag{4.4}$$

Lema 4.1.1. *Los operadores φ y $(G \circ \varphi)$ definidos en (4.2) y (4.4) respectivamente son acotados.*

Demostración:

Sea $F \in C_1^2(\Omega)$, tomemos $\phi = u(F)$ en la ecuación variacional (4.1)

$$\int_0^1 \frac{\partial u}{\partial t} u dx = \int_0^1 \frac{\partial u}{\partial x} \frac{\partial u}{\partial x} dx + \int_0^1 F u dx,$$

integrando respecto a t tenemos

$$\int_0^t \int_0^1 \frac{\partial u}{\partial t} u dx dt = \int_0^t \int_0^1 \left(\frac{\partial u}{\partial x} \right)^2 dx dt + \int_0^t \int_0^1 F u dx dt,$$

ahora

$$\begin{aligned}\int_0^t \int_0^1 \frac{\partial u}{\partial t} u dx dt &= \int_0^1 \left(\int_0^t \frac{\partial u}{\partial t} u dt \right) dx = \int_0^1 \left(\frac{1}{2} \int_0^t \frac{d}{dt} (u)^2 dt \right) dx \\ &= \frac{1}{2} \int_0^1 |u(x, t)|^2 dx - \frac{1}{2} \int_0^1 |u(x, 0)|^2 dx \\ &= \frac{1}{2} \|u\|_{L^2}^2 - \frac{1}{2} \|f\|_{L^2}^2,\end{aligned}$$

4. Fórmulas numéricas

además

$$\frac{1}{2}\|u\|_{L^2}^2 - \frac{1}{2}\|f\|_{L^2}^2 + \int_0^t \int_0^1 \left| \frac{\partial u}{\partial x} \right|^2 dx dt = \int_0^t \int_0^1 F u \, dx \, dt \leq \|F\|_{L^2} \|u\|_{L^2},$$

entonces

$$\begin{aligned} \frac{1}{2}\|u\|_{L^2}^2 - \frac{1}{2}\|f\|_{L^2}^2 &\leq \|F\|_{L^2} \|u\|_{L^2}, \\ \|u\|_{L^2}^2 &\leq \|f\|_{L^2}^2 + 2\|F\|_{L^2} \|u\|_{L^2}. \end{aligned}$$

Por el lema 3.2.4 tenemos que

$$\|u\|_{L^2} \leq \|f\|_{L^2} + 2\|F\|_{L^2}. \quad (4.5)$$

Así φ es acotado.

Puesto que

$$\int_0^t \int_0^1 \left| \frac{\partial u}{\partial x} \right|^2 dx dt \leq \|f\|_{L^2}^2 + 2\|F\|_{L^2} \|u\|_{L^2},$$

y usando (4.5) tenemos que

$$\int_0^t \int_0^1 \left| \frac{\partial u}{\partial x} \right|^2 dx dt \leq \|f\|_{L^2}^2 + 2\|F\|_{L^2} (\|f\|_{L^2} + 2\|F\|_{L^2}),$$

entonces $(G \circ \varphi)$ es acotado.

Lema 4.1.2. *Los operadores φ y $(G \circ \varphi)$ definidos en (4.2) y (4.4) respectivamente son débilmente cerrados.*

Demostración:

Consideremos las siguientes convergencias débiles

$$F_n \rightharpoonup F, \quad \varphi(F_n) \rightharpoonup \varphi^*, \quad G \circ \varphi(F_n) \rightharpoonup G^*,$$

donde $F_n \in C_1^2(\Omega)$. Claramente $F \in C_1^2(\Omega)$, por el lema 4.1.1 $\{\varphi(F_n)\}$ es acotada en $L_2(\Omega)$, entonces existe una subsucesión $\{\varphi(F_{n_k})\}$ que converge débilmente a $u^* \in L_2(\Omega)$ entonces $u^* = u(F)$ donde $\varphi^* = \varphi(F)(x, t) = u(F, x, t)$, ver [7].

Para probar que $(G \circ \varphi)$ es débilmente cerrado, supongamos que

$$G(u_n) \rightharpoonup G^* = (g_o^*, g_1^*).$$

Por definición de G tenemos que

$$\int_{t_1}^{t_2} u_n(F, x, t) dt \rightharpoonup g_o^*, \quad \int_{t_1}^{t_2} \frac{\partial u_n(F, x, t)}{\partial x} dt \rightharpoonup g_1^*.$$

4. Fórmulas numéricas

Como $\varphi(F_n) \rightharpoonup \varphi^*$ y $u(F_n)$ son acotados en $H^1(\Omega)$ aplicamos el teorema de la convergencia dominada para obtener que

$$\int_{t_1}^{t_2} u_n(F, x, t) dt \rightharpoonup \int_{t_1}^{t_2} u(F, x, t) dt = g_o^* \text{ en } L_2(\Omega).$$

Debido a

$$\int_{t_1}^{t_2} \frac{\partial u_n(F, x, t)}{\partial x} dt \rightharpoonup g_1^*,$$

tenemos que para todo $\varphi \in H^1(\Omega)$

$$\lim_{n \rightarrow \infty} \int_0^1 \left(\int_{t_1}^{t_2} \frac{\partial u_n(F, x, t)}{\partial x} dt \right) \varphi dx = \int_0^1 g_1^* \varphi dx.$$

Entonces

$$\begin{aligned} \int_0^1 g_1^* \varphi dx &= \lim_{n \rightarrow \infty} \int_0^1 \left(\int_{t_1}^{t_2} \frac{\partial u_n(F, x, t)}{\partial x} dt \right) \varphi dx \\ &= \lim_{n \rightarrow \infty} \int_{t_1}^{t_2} \left(\int_0^1 \frac{\partial u_n(F, x, t)}{\partial x} \varphi dx \right) dt \\ &= \lim_{n \rightarrow \infty} \int_{t_1}^{t_2} \left(\int_0^1 u_n(F, x, t) \frac{\partial \varphi}{\partial x} dx \right) dt \\ &= \int_{t_1}^{t_2} \left(\int_0^1 u(F, x, t) \frac{\partial \varphi}{\partial x} dx \right) dt \\ &= \int_{t_1}^{t_2} \left(\int_0^1 \frac{\partial u(F, x, t)}{\partial x} \varphi dx \right) dt \\ &= \int_0^1 \left(\int_{t_1}^{t_2} \frac{\partial u(F, x, t)}{\partial x} dt \right) \varphi dx, \end{aligned}$$

para todo $\varphi \in H^1(\Omega)$, entonces

$$g_1^* = \int_{t_1}^{t_2} \frac{\partial u(F, x, t)}{\partial x} (F, x, t) dt,$$

así el operador $(G \circ \varphi)$ es débilmente cerrado.

Para estudiar las soluciones del problema implícito (4.4) definimos el funcional

$$\begin{aligned} J : C_1^2(\Omega) \subset H^1(\Omega) &\longrightarrow \mathbb{R} \\ J(F) &= \|G \circ \varphi(F) - h^\delta\|_Z^2 + \alpha \|F\|_X^2 + \epsilon \|\varphi(F)\|_Y^2, \end{aligned} \quad (4.6)$$

donde

$$h^\delta = \left(\int_{t_1}^{t_2} u^\delta(F, x, t) dt, \int_{t_1}^{t_2} \frac{\partial u^\delta(F, x, t)}{\partial x} dt \right),$$

4. Fórmulas numéricas

supongamos $F^* = 0$, $\varphi^* = 0$ y definamos los espacios $Z = H^1(\Omega) \times L^2[0, 1]$, $X = C_1^2(\Omega)$, $Y = H^1(\Omega)$.

Consideremos el problema de minimización

$$\min_{F \in C_1^2(\Omega)} M_\alpha^\epsilon[G, \varphi, F, h]. \quad (4.7)$$

Para identificar a F se ha creado el siguiente teorema, dado que no se conocía ningún procedimiento para hacerlo a través de esta técnica de regularización.

Teorema 4.1.3. (*Teorema principal*)

La primera variación del funcional J para el problema (4.7) está dado por

$$\delta J(F) = \int_0^1 \int_0^1 (2\alpha F - 2\psi(x, t)) \delta F(u) \, dx \, dt.$$

Para la demostración del teorema primero se introducirán las *ecuaciones de sensibilidad* y la *ecuación del problema adjunto*, que se presentarán en las secciones siguientes. Sin pérdida de generalidad, para el cálculo de las fórmulas numéricas, vamos a suponer que $F(x, t, u) = F(u)$.

4.2. Ecuaciones de sensibilidad

La derivada direccional de J en F , en la dirección $\delta(F)$ es dada por la derivada de Gateaux

$$J'(F, \delta F) = \lim_{s \rightarrow 0} \frac{J(F + s\delta F) - J(F)}{s} = \nabla J(F) \delta F = \delta J. \quad (4.8)$$

Entonces

$$J(F + s\delta F) = \|G \circ \varphi(F + s\delta F) - h^\delta\|_Z^2 + \alpha \|F + s\delta F\|_X^2 + \epsilon \|\varphi(F + s\delta F)\|_Y^2 \quad (4.9)$$

y

$$J(F) = \|G \circ \varphi(F) - h^\delta\|_Z^2 + \alpha \|F\|_X^2 + \epsilon \|\varphi(F)\|_Y^2, \quad (4.10)$$

restando (4.9) y (4.10) tenemos

$$\begin{aligned} J(F + s\delta F) - J(F) &= (\|G \circ \varphi(F + s\delta F) - h^\delta\|_Z^2 - \|G \circ \varphi(F) - h^\delta\|_Z^2) \\ &+ \alpha (\|F + s\delta F\|_X^2 - \|F\|_X^2) + \epsilon (\|\varphi(F + s\delta F)\|_Y^2 - \|\varphi(F)\|_Y^2), \end{aligned}$$

desarrollando cada paréntesis, dividiendo sobre s y calculando el límite cuando $s \rightarrow 0$ tenemos

4. Fórmulas numéricas

$$\begin{aligned}
1. & (\|G \circ \varphi(F + s\delta F) - h^\delta\|_Z^2 - \|G \circ \varphi(F) - h^\delta\|_Z^2) \\
&= \int_0^1 \int_0^1 \sum_{i=1}^n ((u(x, t, F + s\delta F) - h^\delta(x, t))^2 - (u(x, t, F) - h^\delta(x, t))^2) \sigma(x - x_i) dx dt \\
&= \int_0^1 \int_0^1 \sum_{i=1}^n \frac{1}{s} (u(x, t, F + s\delta F)^2 - u(x, t, F)^2 \\
&\quad - 2h^\delta(x, t)(u(x, t, F + s\delta F) - u(x, t, F))) \sigma(x - x_i) dx dt \\
&= 2 \int_0^1 \int_0^1 \sum_{i=1}^n (u(x, t, F) - h^\delta(x, t)) \sigma(x - x_i) \delta u dx dt,
\end{aligned}$$

donde σ es la función de Dirac.

$$\begin{aligned}
2. & (\|F + s\delta F\|_X^2 - \|F\|_X^2) \\
&= \langle F + s\delta F, F + s\delta F \rangle - \langle F, F \rangle \\
&= 2s \langle F, \delta F \rangle + s^2 \langle \delta F, \delta F \rangle = \lim_{s \rightarrow 0} (\frac{1}{s} 2s \langle F, \delta F \rangle + \frac{1}{s} s^2 \langle \delta F, \delta F \rangle) \\
&= \lim_{s \rightarrow 0} (2 \langle F, \delta F \rangle + s \langle \delta F, \delta F \rangle) = 2 \langle F, \delta F \rangle = 2 \int_0^1 \int_0^1 F \delta F dx dt. \\
3. & (\|\varphi(F + s\delta F)\|_Y^2 - \|\varphi(F)\|_Y^2) \\
&= \langle \varphi(F + s\delta F), \varphi(F + s\delta F) \rangle - \langle \varphi(F), \varphi(F) \rangle \\
&= \langle \varphi(F + s\delta F) - \varphi(F) + \varphi(F), \varphi(F + s\delta F) \rangle - \langle \varphi(F), \varphi(F) \rangle \\
&= \langle \varphi(F + s\delta F) - \varphi(F), \varphi(F + s\delta F) \rangle + \langle \varphi(F), \varphi(F + s\delta F) \rangle - \langle \varphi(F), \varphi(F) \rangle \\
&= \langle \varphi(F + s\delta F) - \varphi(F), \varphi(F + tv) \rangle + \langle \varphi(F), \varphi(F + s\delta F) - \varphi(F) \rangle \\
&= \lim_{s \rightarrow 0} \frac{1}{s} (\langle (\varphi(F + s\delta F) - \varphi(F)), \varphi(F + s\delta F) \rangle + \langle \varphi(F), (\varphi(F + s\delta F) - \varphi(F)) \rangle) \\
&= \langle \delta(\varphi(F)), \varphi(F) \rangle + \langle \varphi(F), \delta(\varphi(F)) \rangle = 2 \langle \delta(\varphi(F)), \varphi(F) \rangle \\
&= 2 \int_0^1 \int_0^1 u \delta u dx dt.
\end{aligned}$$

Finalmente tenemos que

$$J'(F, \delta F) = 2\alpha \int_0^1 \int_0^1 F \delta F dx dt + \int_0^1 \int_0^1 (2 \sum_{i=1}^n (u - h^\delta) \sigma(x - x_i) + 2\epsilon u) \delta u dx dt.$$

Observemos que en el cálculo de $J'(F)$ aparece el término δu en la segunda integral, la idea es eliminarla y para esto primero introducimos las ecuaciones de sensibilidad.

$$\frac{\partial u}{\partial t}(x, t, F + s\delta F) = \frac{\partial^2 u}{\partial x^2}(x, t, F + s\delta F) + [F + s\delta F](u(x, t, F + s\delta F)) \quad (4.11)$$

y

$$\frac{\partial u}{\partial t}(x, t, F) = \frac{\partial^2 u}{\partial x^2}(x, t, F) + F(u(x, t, F)), \quad (4.12)$$

restando (4.11) y (4.12) tenemos

$$\frac{\partial u}{\partial t}(x, t, F + s\delta F) - \frac{\partial u}{\partial t}(x, t, F)$$

4. Fórmulas numéricas

$$= \frac{\partial^2}{\partial x^2} [u(x, t, F + s\delta F) - u(x, t, F)] + [F + s\delta F](u(x, t, F + s\delta F)) - F(u(x, t, F)),$$

dividiendo por s y calculando el límite cuando $s \rightarrow 0$ tenemos

$$\begin{aligned} \frac{\partial}{\partial t}(\delta u) &= \frac{\partial^2}{\partial x^2}(\delta u) + \lim_{s \rightarrow 0} F[u(x, t, F + s\delta F)] + s\delta F[u(x, t, F + s\delta F)] - F[u(x, t, F)] \\ &= \frac{\partial^2}{\partial x^2}(\delta u) + \delta F[u(x, t, F)] + \delta F[u(x, t, F)]. \end{aligned}$$

Así, tenemos las siguientes **Ecuaciones de sensibilidad**

$$\begin{aligned} \frac{\partial}{\partial t}(\delta u) &= \frac{\partial^2}{\partial x^2}(\delta u) + 2\delta F(u), \\ \delta u(0, t, F) &= 0, \\ \delta u(1, t, F) &= 0, \\ \delta u(x, 0, F) &= 0. \end{aligned} \tag{4.13}$$

Las condiciones iniciales y de frontera se determinan de forma análoga. Ahora como ya encontramos las ecuaciones de sensibilidad, a continuación vamos a encontrar la ecuación diferencial del problema adjunto al problema inicial (1)-(2).

4.3. Ecuación del problema adjunto

Para introducir la ecuación diferencial del operador adjunto, utilizamos el concepto de multiplicadores de Lagrange generalizados. Con este objetivo en mente, consideremos

$$\begin{aligned} \Gamma &= J + \int_0^1 \int_0^1 \psi(x, t)[-u_{xx} + u_t - F(u)]dxdt \\ &\quad + \int_0^1 \theta(t)[u(1, t) - j(t)]dt + \int_0^1 \eta(t)[u(0, t) - g(t)]dt + \int_0^1 \theta(x)[u(x, 0) - f(x)]dx, \end{aligned}$$

donde $\psi : \Omega \rightarrow \mathbb{R}$, $\theta, \eta, : [0, 1] \rightarrow \mathbb{R}$ son multiplicadores de Lagrange generalizados.

Calculemos $\delta\Gamma$:

$$\delta\Gamma = \lim_{s \rightarrow 0} \frac{\Gamma(F + s\delta F) - \Gamma(F)}{s}. \tag{4.14}$$

4. Fórmulas numéricas

Sea

$$\begin{aligned}\Gamma(F + s\delta F) &= J(F + s\delta F) + \int_0^1 \int_0^1 \psi(x, t)[-u_{xx}(x, t, F + s\delta F) + u_t(x, t, F + s\delta F) \\ &\quad - [F + s\delta F](u(x, t, s + \delta F))]dxdt + \int_0^1 \theta(t)[u(1, t, F + s\delta F) - j(t)]dt \\ &\quad + \int_0^1 \eta(t)[u(0, t, F + s\delta F) - g(t)]dt \\ &\quad + \int_0^1 \theta(x)[u(x, 0, F + s\delta F) - f(x)]dx\end{aligned}$$

y

$$\begin{aligned}\Gamma(F) &= J(F) + \int_0^1 \int_0^1 \psi(x, t)[-u_{xx}(x, t, F) + u_t(x, t, F) - [F](u(x, t, F))]dxdt \\ &\quad + \int_0^1 \theta(t)[u(1, t, F) - j(t)]dt + \int_0^1 \eta(t)[u(0, t, F) - g(t)]dt \\ &\quad + \int_0^1 \theta(x)[u(x, 0, F) - f(x)]dx,\end{aligned}$$

ahora

$$\begin{aligned}\Gamma(F + s\delta F) - \Gamma(F) &= (J(F + s\delta F) - J(F)) \\ &\quad + \int_0^1 \int_0^1 \psi(x, t)[-u_{xx}(x, t, F + s\delta F) + u_t(x, t, F + s\delta F) \\ &\quad - [F + s\delta F](u(x, t, s + \delta F)) \\ &\quad + u_{xx}(x, t, F) - u_t(x, t, F) + [F](u(x, t, F))]dxdt \\ &\quad + \int_0^1 \theta(t)[u(1, t, F + s\delta F) - u(1, t, F)]dt \\ &\quad + \int_0^1 \eta(t)[u(0, t, F + s\delta F) - u(0, t, F)]dt \\ &\quad + \int_0^1 \theta(x)[u(x, 0, F + s\delta F) - u(x, 0, F)]dx.\end{aligned}$$

Dividiendo sobre s , calculando el límite cuando $s \rightarrow 0$ y haciendo uso de las ecuaciones de sensibilidad (4.13) tenemos

$$\delta\Gamma = \delta J + \int_0^1 \int_0^1 \psi(x, t)\left[\frac{\partial}{\partial t}(\delta u) - \frac{\partial^2}{\partial x^2}(\delta u) - 2\delta F(u)\right]dxdt.$$

Ahora tomemos

4. Fórmulas numéricas

$$\begin{aligned} I_1 &= - \int_0^1 \int_0^1 \psi(x, t) \frac{\partial^2}{\partial x^2} (\delta u) dx dt, \\ I_2 &= - \int_0^1 \int_0^1 \psi(x, t) 2\delta F(u) dx dt, \\ I_3 &= \int_0^1 \int_0^1 \psi(x, t) \frac{\partial}{\partial t} (\delta(u)) dx dt. \end{aligned}$$

Integrando por partes tenemos

$$\begin{aligned} I_1 &= - \int_0^1 \int_0^1 \psi_{xx} \delta u dx dt, \\ I_2 &= - \int_0^1 \int_0^1 \psi(x, t) 2\delta F(u) dx dt, \\ I_3 &= - \int_0^1 \int_0^1 \psi_t \delta(u) dx dt. \end{aligned}$$

Puesto que

$$\psi(1, t) = 0, \quad \psi(0, t) = 0, \quad \psi_x(1, 0) = 0, \quad \psi(x, 1) = 0.$$

Entonces

$$\begin{aligned} \delta\Gamma &= 2\alpha \int_0^1 \int_0^1 F \delta F(u) dx dt + \int_0^1 \int_0^1 (2 \sum_{i=1}^n (u - h^\delta) \sigma(x - x_i) + 2\epsilon u) \delta u dx dt \\ &\quad - \int_0^1 \int_0^1 \psi_{xx} \delta u dx dt - \int_0^1 \int_0^1 \psi(x, t) 2\delta F(u) dx dt - \int_0^1 \int_0^1 \psi_t \delta u dx dt, \end{aligned}$$

agrupando términos tenemos:

$$\begin{aligned} \delta\Gamma &= \int_0^1 \int_0^1 (2 \sum_{i=1}^n (u - h^\delta) \sigma(x - x_i) + 2\epsilon u - \psi_{xx} - \psi_t) \delta u dx dt \\ &\quad + \int_0^1 \int_0^1 (2\alpha F + 2\psi(x, t)) \delta F(u) dx dt. \end{aligned}$$

De esta forma, la **ecuación del operador adjunto** está dada por

$$\begin{aligned} \psi_t &= -\psi_{xx} + 2 \sum_{i=1}^n (u - j^\delta) \sigma(x - x_i) + 2\epsilon u, \\ \psi(1, t) &= \psi(0, t) = 0, \\ \psi_x(1, 0) &= \psi(x, 1) = 0. \end{aligned} \tag{4.16}$$

4.4. Gradiente del funcional

Demostración: Una vez resuelta la ecuación del adjunto (4.16), demostramos que la ecuación (4.15) queda de la forma

$$\delta J(F) = \int_0^1 \int_0^1 (2\alpha F - 2\psi(x, t)) \delta F(u) dx dt,$$

lo cual prueba el teorema 4.1.3.

Ahora con el fin de estudiar numéricamente el problema (3.3), usamos splines lineales para aproximar la función F :

$$F(u(x, t)) = \sum_{i=1}^N F_i \phi_i(x) \phi_i(t),$$

donde

$$\phi_i(x) = \begin{cases} 1 & \text{si } x \in [x_i, x_{i+1}] \\ 0 & \text{si } x \text{ en otro caso} \end{cases},$$

por lo tanto el gradiente del funcional es

$$\frac{\partial J}{\partial F_i} = \int_0^1 \int_0^1 -2\psi \phi_i(x) \phi_i(t) dx dt + 2\alpha F_i. \quad (4.17)$$

En el proceso de minimizar J se utiliza el método del gradiente conjugado, de la siguiente manera.

Sea $\mathbf{p} = [F_0, F_1, \dots, F_N]^T \in \mathbb{R}^{N+1}$, los parámetros sucesivos obtenidos por

$$\mathbf{p}^{s+1} = \mathbf{p}^s + \gamma^s \mathbf{d}^s, \quad s = 0, 1, 2, \dots,$$

donde s es el número de iteración, γ^s es el parámetro de paso del descenso, $\mathbf{d}^s$ es la dirección de descenso y $\mathbf{p}^0$ es el paso inicial para el vector no conocido $\mathbf{p}$.

La dirección de descenso es calculada por

$$\mathbf{d}^s = -\mathbf{g}^s + \beta^s \mathbf{d}^{s-1},$$

donde $\mathbf{g}$ es el gradiente del funcional que se minimizará y β^s es definido como

$$\beta^0 = 0, \quad \beta^s = \frac{\langle \mathbf{g}^s - \mathbf{g}^{s-1}, \mathbf{g}^s \rangle}{\|\mathbf{g}^{s-1}\|^2},$$

donde $\langle \cdot, \cdot \rangle$ y $\|\cdot\|$ denota el producto interior y la norma en $\mathbb{R}^{N+1}$ respectivamente.

4. Fórmulas numéricas

4.5. Mejoramiento de la estimación de F

Finalmente para el mejorar la estimación de F , debemos encontrar la solución del problema de minimización en una dimensión. Estos cálculos los realizaremos para la función $F(x, t)$.

$$\gamma^s = \min_{\gamma > 0} J(\mathbf{p}^s - \gamma^s \mathbf{d}^s).$$

Usamos la siguiente la notación $\mathbf{p}^s \equiv F$ y $\gamma^s \mathbf{d}^s \equiv \gamma(\nabla u, \nabla J)$, y dado que $u(x, t, F - \gamma \nabla J(F)) \simeq u(x, t, F) - \gamma(\nabla u, \nabla J)$, tenemos la aproximación

$$\begin{aligned} J &\simeq \int_0^1 \int_0^1 \sum_{i=1}^n (u - \gamma(\nabla u, \nabla J) - h^\delta)^2 \sigma(x - x_i) dx dt + \epsilon \int_0^1 \int_0^1 (u - \gamma(\nabla u, \nabla J))^2 dx dt \\ &\quad + \alpha \sum_{i=0}^N (F_i - \gamma \frac{\partial J}{\partial F_i})^2 \end{aligned}$$

y

$$\begin{aligned} \frac{dJ}{d\gamma} &= -2 \int_0^1 \int_0^1 \sum_{i=1}^n (u - \gamma(\nabla u, \nabla J) - h^\delta) \sigma(x - x_i) (\nabla u, \nabla J) dx dt \\ &\quad - 2\epsilon \int_0^1 \int_0^1 (u - \gamma(\nabla u, \nabla J)) (\nabla u, \nabla J) dx dt - 2\alpha \sum_{i=0}^N (F_i - \gamma \frac{\partial J}{\partial F_i}) (\frac{\partial J}{\partial F_i}). \end{aligned}$$

Haciendo $\frac{dJ}{d\gamma} = 0$ y solucionando para γ tenemos

$$\gamma = \frac{\int_0^1 \int_0^1 \sum_{i=1}^n (u - g^\delta) (\nabla u, \nabla J) \sigma(x - x_i) dx dt + \epsilon \int_0^1 \int_0^1 u (\nabla u, \nabla J) dx dt + \alpha \sum_{i=0}^N F_i \frac{\partial J}{\partial F_i}}{(1 + \epsilon) \int_0^1 \int_0^1 (\nabla u, \nabla J)^2 dx dt + \alpha \sum_{i=0}^N \left(\frac{\partial J}{\partial F_i} \right)^2}.$$

Podemos conocer $\nabla u = \frac{\partial u}{\partial F_i}$, que son las componentes de la variable $\delta u = \frac{\partial u}{\partial F_i} \delta F_i$ y estas pueden ser encontradas solucionando las ecuaciones de sensibilidad para $\theta = \theta(x, t, \mathbf{d}^s)$. En el método del gradiente conjugado ∇J se relaciona con $\mathbf{d}^s$ y haciendo uso de la derivada de Gateaux para encontrar una expresión para δu tenemos

$$\delta u(F_n) = \lim_{v \rightarrow 0} \frac{u(F_n + v \mathbf{d}^s) - u(F_n)}{v} = (\nabla u, \mathbf{d}^s) = (\nabla u, \nabla J).$$

Al sustituir δu por $\theta(x, t, \mathbf{d}^s)$ en la ecuación de sensibilidad (4.7), tenemos

$$\frac{\partial \theta}{\partial t} = \frac{\partial^2 \theta}{\partial x^2} + 2\delta(F(x, t)),$$

4. Fórmulas numéricas

reemplazando δF por $\sum_{i=0}^N d_i \phi_i(x) \phi_i(t)$ tenemos

$$\frac{\partial \theta}{\partial t} = \frac{\partial^2 \theta}{\partial x^2} + 2 \sum_{i=0}^N d_i \phi_i(x) \phi_i(t).$$

En conclusión el parámetro de descenso queda de la forma

$$\gamma^s = \frac{\int_0^1 \int_0^1 \sum_{i=1}^n (u - h^\delta) \sigma(x - x_i) \theta(x, t; \mathbf{d}^s) dx dt + \epsilon \int_0^1 \int_0^1 \theta(x, t; \mathbf{d}^s) dx dt + \alpha \sum_{i=0}^N F_i^s d_i^s}{(1 + \epsilon) \int_0^1 \int_0^1 (\theta(x, t; \mathbf{d}^s))^2 dx dt + \alpha \sum_{i=0}^N (d_i^s)^2},$$

donde $\theta(x, t; \mathbf{d}^s)$ es la solución del problema

$$\begin{aligned} \frac{\partial \theta}{\partial t} &= \frac{\partial^2 \theta}{\partial x^2} + 2 \sum_{i=0}^N d_i \phi_i(x) \phi_i(t), \\ \theta(0, t) &= 0, \\ \theta(1, t) &= 0, \\ \theta(x, 0) &= 0. \end{aligned}$$

4.6. Experimentación numérica

Podemos ver que si u no depende del tiempo en la ecuación $u_t(x, t) = u_{xx}(x, t) + F(x, t, u)$, esta ecuación se convierte en un problema elíptico y es el caso más sencillo de resolver numéricamente. Los ejemplos numéricos los realizaremos sobre la ecuación

$$-u_{xx} = F(u),$$

y las fórmulas numéricas para esta ecuación se deducen de los cálculos numéricos del problema (1)-(2).

- Ecuación del adjunto

$$\begin{aligned} 2 \sum_{i=1}^n (u - j^\delta) \sigma(x - x_i) + 2\epsilon u - \psi_{xx} &= 0, \\ \psi(1) = \psi(0) = \psi_x(1) &= 0. \end{aligned}$$

- Gradiente del funcional

$$\frac{\partial J}{\partial F_i} = \int_0^1 -2\psi \phi_i(x) dx + 2\alpha F_i.$$

4. Fórmulas numéricas

- Parámetro de descenso

$$\gamma^s = \frac{\int_0^1 \sum_{i=1}^n (u - h^\delta) \sigma(x - x_i) \theta(x; \mathbf{d}^s) dx + \epsilon \int_0^1 \theta(x; \mathbf{d}^s) dx + \alpha \sum_{i=0}^N F_i^s d_i^s}{(1 + \epsilon) \int_0^1 (\theta(x; \mathbf{d}^s))^2 dx + \alpha \sum_{i=0}^N (d_i^s)^2},$$

donde θ es la solución del problema

$$\begin{aligned} -\frac{\partial^2 \theta}{\partial x^2} &= 2 \sum_{i=0}^N d_i \phi_i(x), \\ \theta(0) &= \theta(1) = 0. \end{aligned}$$

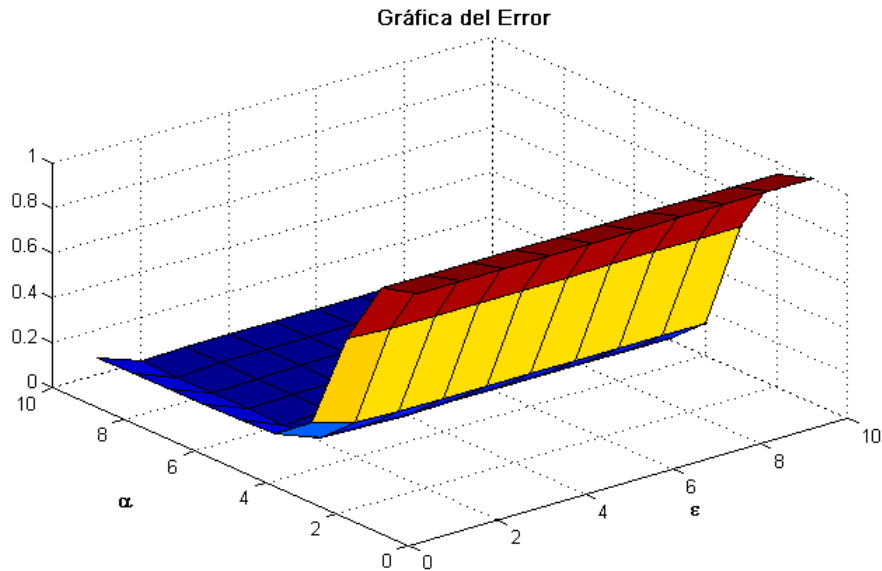
Los datos u^δ son generados agregando números aleatorios distribuidos uniformemente a la solución exacta u . El programa que usamos para nuestros ejemplos numéricos es MATLAB 7.8.0 (R2009a), computador Asus, con procesador intel CORE i3.

Ejemplo 1

$$\begin{aligned} -u_{xx}(x) &= F(u), \\ u(0) &= u(1) = 0, \end{aligned}$$

donde las soluciones exactas son $u(x) = \text{sen}(\pi x)$ con $F(u) = \pi^2 u$, con error en los datos $u^\delta = u + \delta$ ($\delta = 0,02$), el proceso que llevaremos a cabo en la identificación de F es el siguiente

Primero identificaremos $H(x) = F(u(x))$, en este proceso identificamos simultáneamente $u(x)$, luego conocidos u y H determinaremos F .



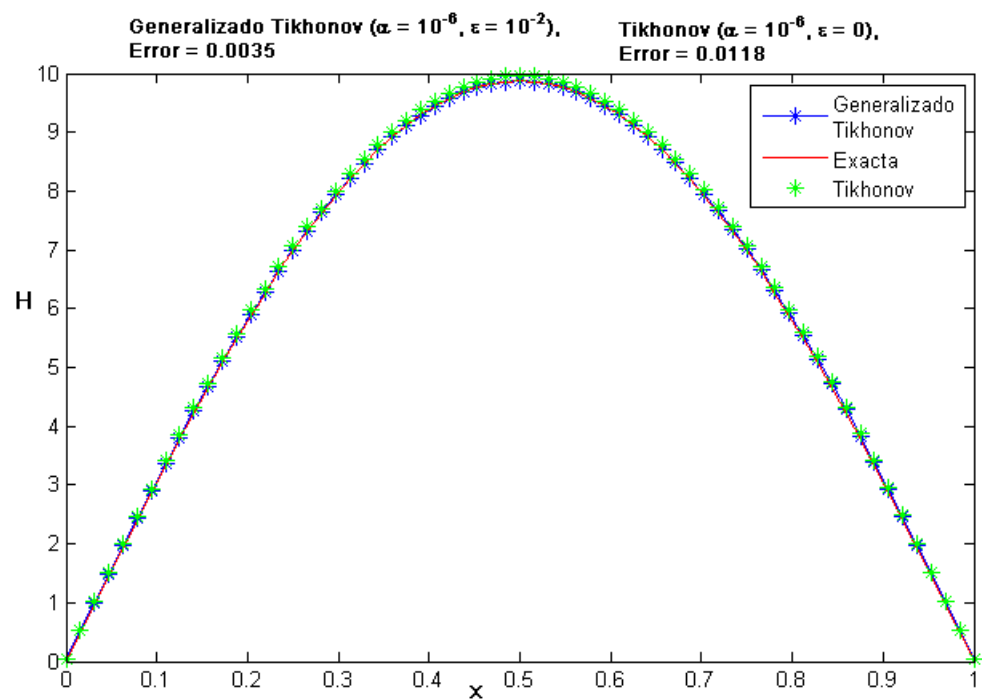
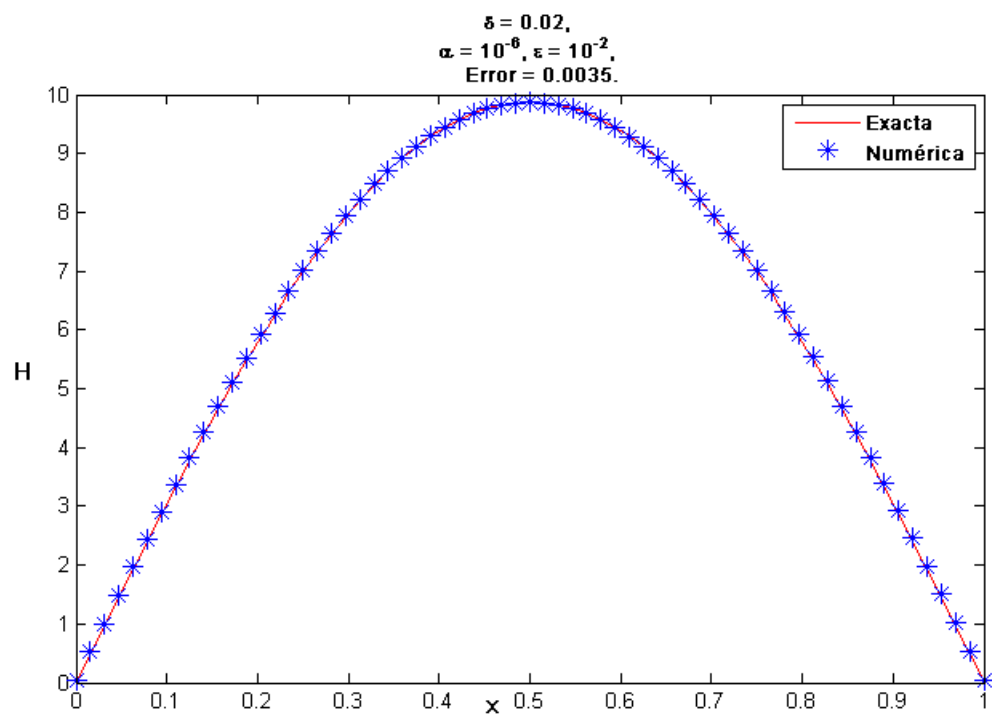
4. Fórmulas numéricas

Tabla de errores

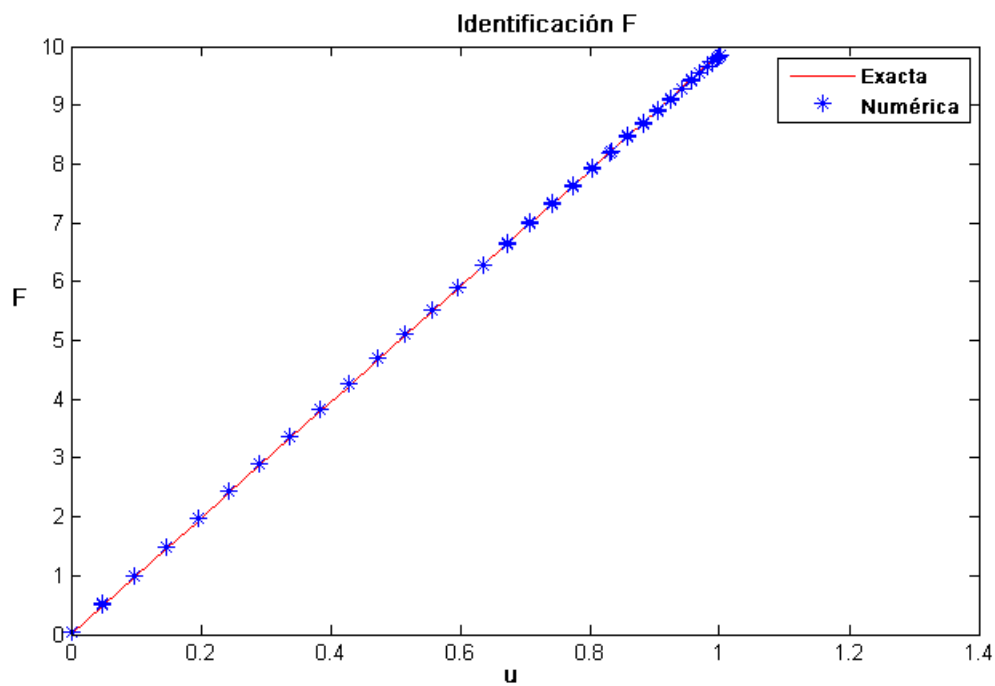
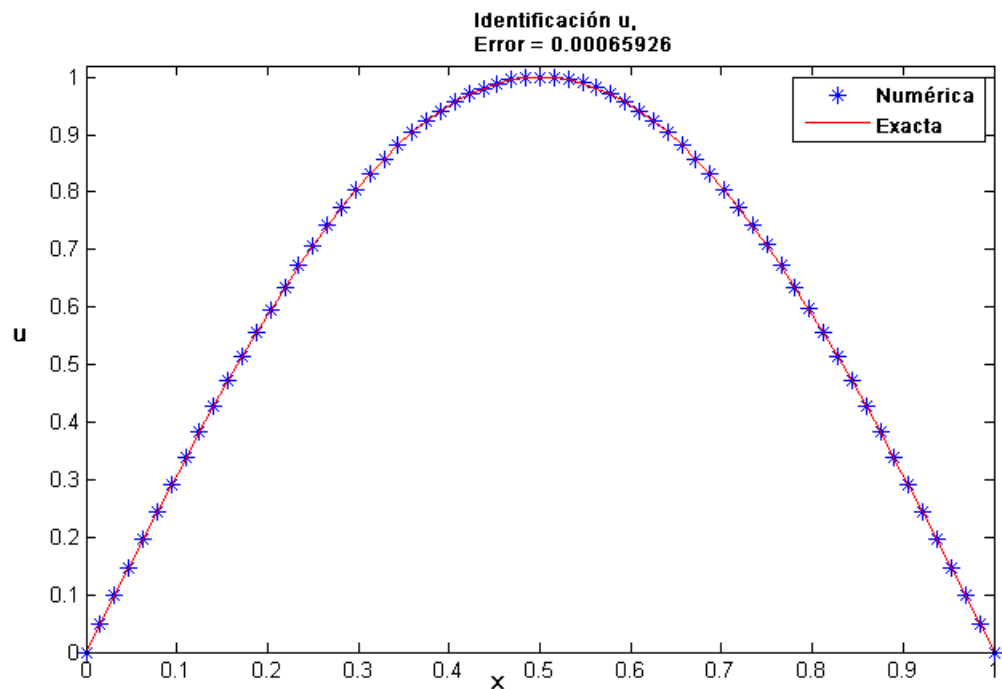
$\alpha \backslash \epsilon$	10^{-1}	10^{-2}	10^{-3}	10^{-4}	10^{-5}	10^{-6}	10^{-7}	10^{-8}	10^{-9}	10^{-10}
10^{-1}	0.9945	0.9945	0.9945	0.9945	0.9945	0.9945	0.9945	0.9945	0.9945	0.9945
10^{-2}	0.9479	0.9476	0.9476	0.9476	0.9476	0.9476	0.9476	0.9476	0.9476	0.9476
10^{-3}	0.6545	0.6436	0.6423	0.6423	0.6423	0.6423	0.6423	0.6423	0.6423	0.6423
10^{-4}	0.2097	0.1501	0.1430	0.1429	0.1429	0.1429	0.1429	0.1429	0.1429	0.1429
10^{-5}	0.0929	0.0136	0.0049	0.0041	0.0040	0.0040	0.0040	0.0040	0.0040	0.0040
10^{-6}	0.0797	0.0035	0.0114	0.0123	0.0124	0.0124	0.0124	0.0124	0.0124	0.0124
10^{-7}	0.0781	0.0040	0.0127	0.0136	0.0137	0.0137	0.0137	0.0137	0.0137	0.0137
10^{-8}	0.0782	0.0045	0.0129	0.0138	0.0139	0.0139	0.0139	0.0139	0.0139	0.0139
10^{-9}	0.0783	0.0049	0.0131	0.0140	0.0141	0.0141	0.0141	0.0141	0.0141	0.0141
10^{-10}	0.0784	0.0054	0.0133	0.0144	0.0149	0.0149	0.0149	0.0149	0.0149	0.0149

Como podemos observar los valores de parámetros para los cuales se tiene el menor error son: $\alpha = 10^{-6}$ y $\epsilon = 10^{-2}$.

4. Fórmulas numéricas



4. Fórmulas numéricas



Como podemos observar F es una función lineal de la forma $F(u) = cu$ donde debemos determinar el valor de c , numéricamente haciendo F/u tenemos que $c \approx \pi^2$.

4. Fórmulas numéricas

Ejemplo 2

$$-u_{xx}(x) = F(u),$$

$$u(0) = u(1) = 0,$$

donde las soluciones exactas son $u(x) = \cos(\frac{\pi}{2}(2x + 1))$ con $F(u) = \pi^2 u$, con error en los datos u^δ ($\delta = 0,02$).

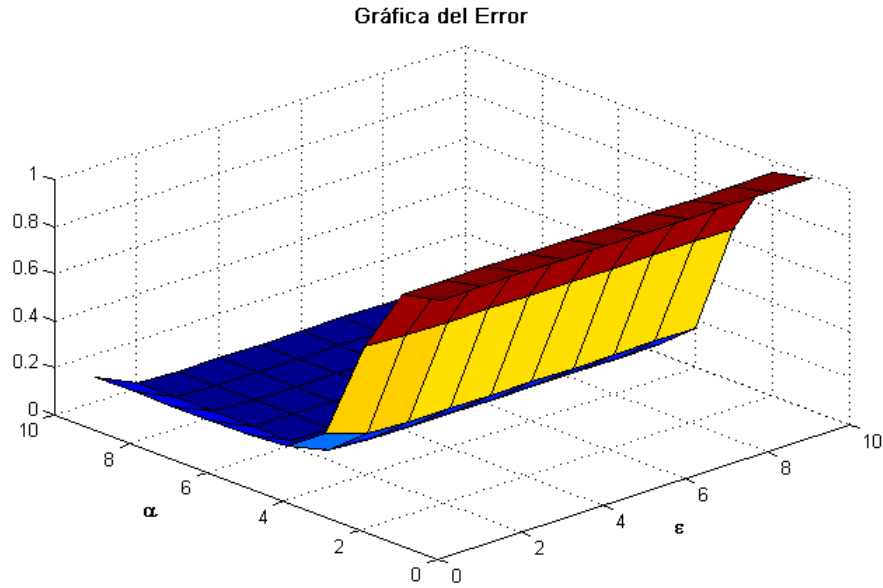
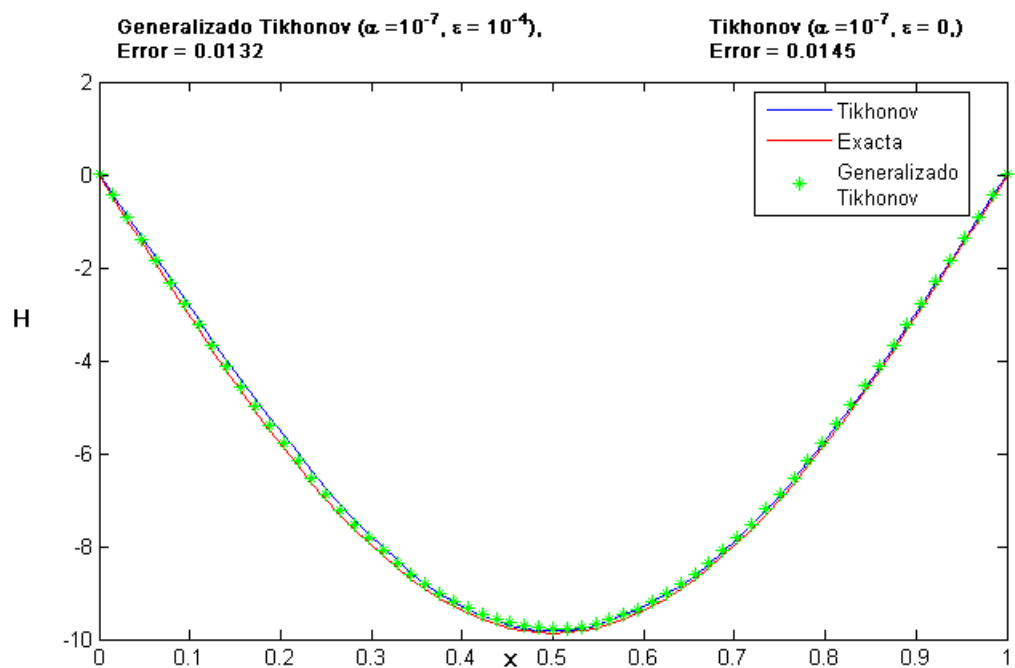
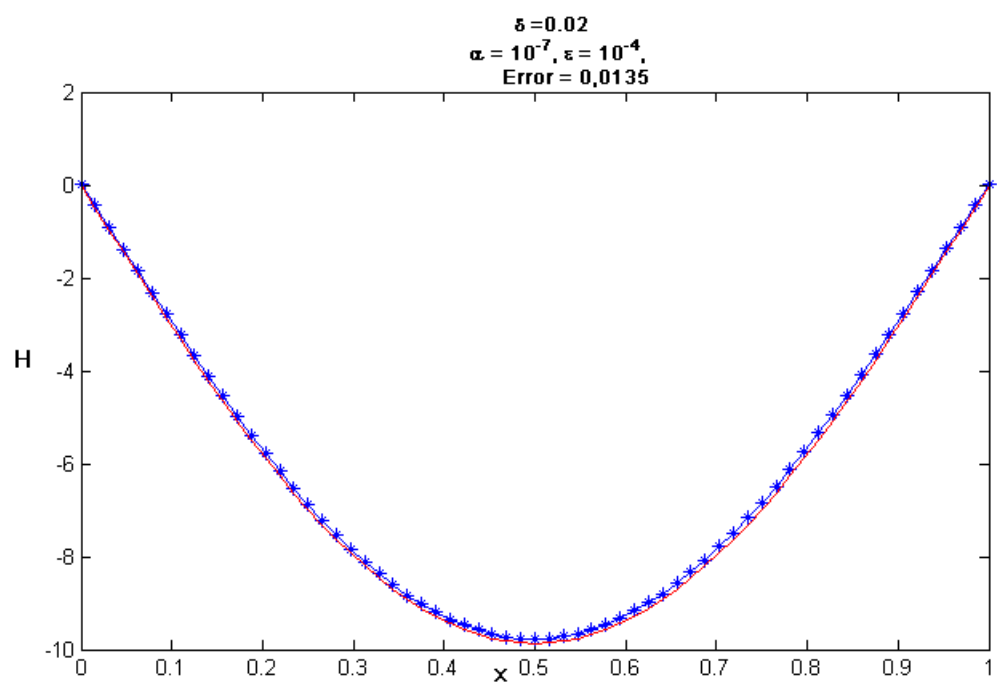


Tabla de errores

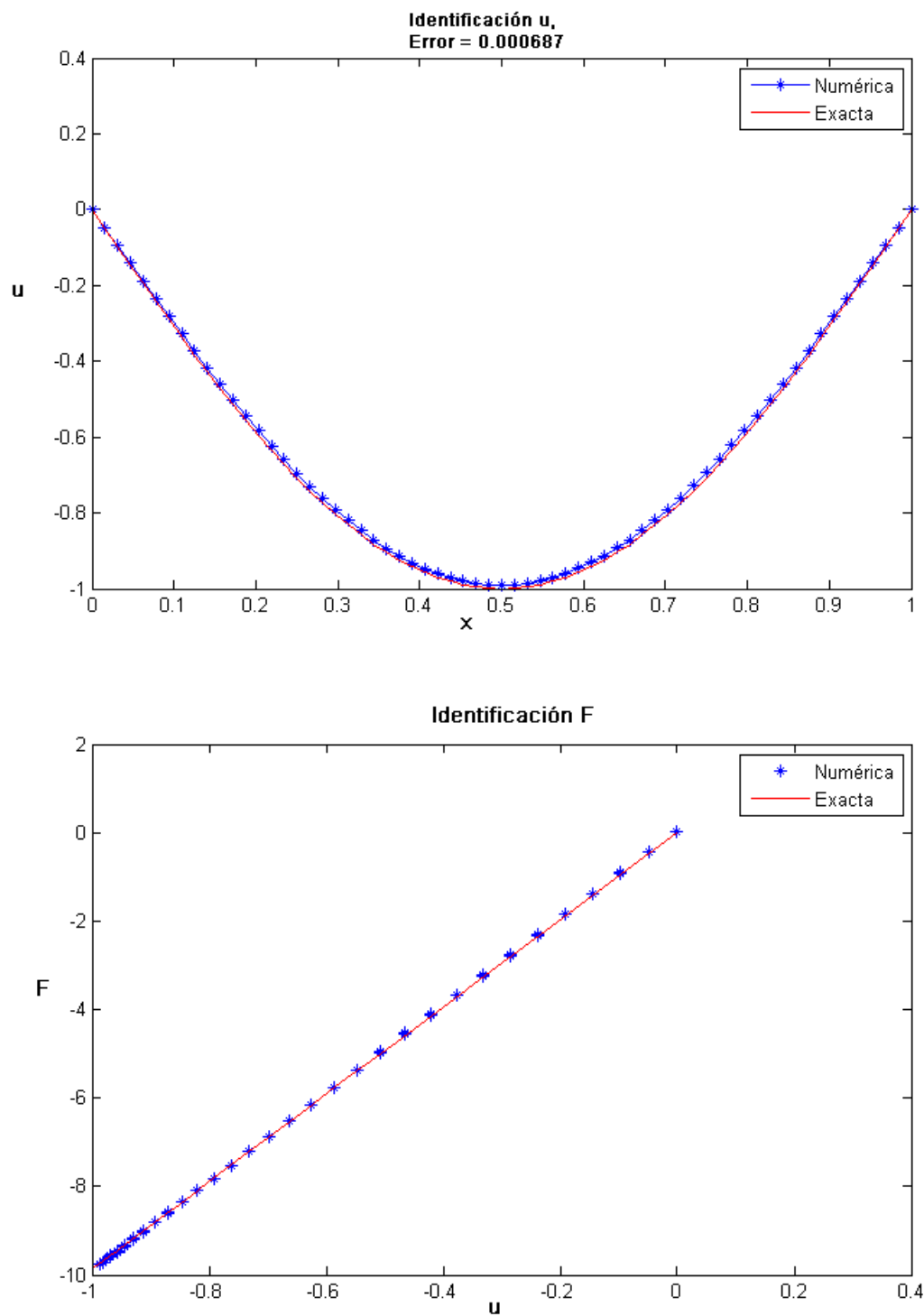
$\alpha \backslash \epsilon$	10^{-1}	10^{-2}	10^{-3}	10^{-4}	10^{-5}	10^{-6}	10^{-7}	10^{-8}	10^{-9}	10^{-10}
10^{-1}	0.9861	0.9861	0.9861	0.9861	0.9861	0.9861	0.9861	0.9861	0.9861	0.9861
10^{-2}	0.9482	0.9479	0.9479	0.9479	0.9479	0.9479	0.9479	0.9479	0.9479	0.9479
10^{-3}	0.6632	0.6526	0.6515	0.6514	0.6514	0.6514	0.6514	0.6514	0.6514	0.6514
10^{-4}	0.2302	0.1722	0.1659	0.1653	0.1652	0.1652	0.1652	0.1652	0.1652	0.1652
10^{-5}	0.1166	0.0393	0.0309	0.0300	0.0299	0.0299	0.0299	0.0299	0.0299	0.0299
10^{-6}	0.1017	0.0238	0.0153	0.0144	0.0143	0.0143	0.0143	0.0143	0.0143	0.0143
10^{-7}	0.1017	0.0226	0.0143	0.0135	0.0135	0.0135	0.0135	0.0135	0.0135	0.0135
10^{-8}	0.1017	0.0226	0.0145	0.0138	0.0137	0.0137	0.0137	0.0137	0.0137	0.0137
10^{-9}	0.1017	0.0226	0.0146	0.0139	0.0138	0.0138	0.0138	0.0138	0.0138	0.0138
10^{-10}	0.1017	0.0227	0.0146	0.0139	0.0138	0.0138	0.0138	0.0138	0.0138	0.0138

Como podemos observar los valores de parámetros para los cuales se tiene el menor error son: $\alpha = 10^{-7}$ y $\epsilon = 10^{-4}$.

4. Fórmulas numéricas



4. Fórmulas numéricas



Como podemos observar F es una función lineal de la forma $F(u) = cu$ donde debemos determinar el valor de c , numéricamente haciendo F/u tenemos que $c \approx \pi^2$.

Capítulo 5

Conclusiones

Durante el presente trabajo vimos la importancia de F en la ecuación $u_t(x, t) = u_{xx}(x, t) + F(x, t, u)$ en las diferentes aplicaciones en la ciencia y la ingeniería, de ahí nuestro interés en plantearnos el problema inverso de identificación de coeficientes a nuestro problema, es decir, identificar F . La importancia de este trabajo en el proceso de identificación de F , es debido a que usamos un método de regularización reciente [7], llamado generalizado de Tikhonov que consiste de dos parámetros α y ϵ llamados parámetros de regularización y donde definimos dos operadores G y φ adecuados que cumplen con las condiciones del método, así, definimos el problema de identificación de F como un problema inverso implícito

$$G(\varphi(F)) = \left(\int_{t_1}^{t_2} u(F, x, t) dt, \int_{t_1}^{t_2} \frac{\partial u(F, x, t)}{\partial x} dt \right).$$

El método generalizado de Tikhonov consiste en definir el funcional J y paso siguiente resolver el problema de minimización (P_ϵ^α) , dando así, la primera variación del funcional, δJ , (Teorema principal 4.1.3) a través de las ecuaciones de sensibilidad y la ecuación del problema del adjunto. Por otro lado, los ejemplos numéricos presentados, se realizaron para el caso donde u no depende de t y bajo ciertas condiciones sobre F donde ella es monótona no decreciente en u , queremos resalta que este problema es más sencillo de resolver numéricamente más no fácil y para estos ejemplos podemos concluir lo siguiente

- El método generalizado de Tikhonov, nos proporciona mejores aproximaciones a la exacta que el método de Tikhonov (un sólo parámetro), aunque la diferencia del error no es muy grande, el hecho de tener dos parámetros de regularización (α y ϵ) brinda más posibilidades de hacer mejores aproximaciones.
- Para nuestros ejemplos numéricos trabajamos con errores en los datos ($\delta = 0,02$) y el método aplica una buena regularización para encontrar la aproximación a F , en artículos mencionados anteriormente no trabajan con errores en los datos.
- Por último, en el proceso de identificación de F identificamos simultáneamente u y esta aproximación tiene un error más pequeño comparado con el error de F .

5. Conclusiones

Trabajo futuro

Presentar ejemplos numéricos donde se aplican las fórmulas numéricas desarrolladas en el capítulo 4, para la identificación de F en los siguientes problemas

■

$$u_t(x, t) = u_{xx}(x, t) + F(x, t, u), \quad 0 \leq x \leq 1, \quad 0 \leq t \leq 1.$$

■

$$u_t(x, t) = u_{xx}(x, t) + F(u), \quad 0 \leq x \leq 1, \quad 0 \leq t \leq 1.$$

Capítulo 6

Referencias

- [1] Jorge Nocedal, Stephen J. Wright, *Numerical Optimization*, Springer, 1999.
- [2] Cannon, J.R and DuChateau, P. *Structural identification of an unknown source term in a heat equation*, Inverse problems, 14, 535-551. 1998.
- [3] Pilant, M. and Rundell, W. *Undetermined coefficient problems for quasilinear parabolic equations*, proceedings of the conference on inverse problems in partial differential equation, Arcata, California, July 29- august 4. 1989.
- [4] Engl-Hanke-Neubauer, *Regularization of inverse problems*, 1996.
- [5] Afet Golayoglu Fatullayev, *Numerical Method of identification of an Unknown Source term in a heat equation*, Mathematical problems in engineering, 2002.
- [6] George E.Forsythe, *Finite-Difference Methods for Partial Differential Equations*, 1960.
- [7] Hinestroza, D., Murio D. A., Shenghe Z, *Regularization Techniques for nonlinear problems.*, Journal of Comput. Math. Applic. 1999.
- [8] Wei Chenga, Ling-Ling Zhaob, Chu-Li Fuc, *Source term identification for an axisymmetric inverse heat conduction problem*. Computers and Mathematics with Applications, 2009.
- [9] Diego A. Murio, *Mollification and space marching*, 2003.
- [10] Mejía C.E. and D.A. Murio, *Mollified Hyperbolic Method for Coefficient Identification Problems*, Computers Math. Applic., (1993).
- [11] C.V. Pao *Nonlinear Parabolic and Elliptic equations*, 1933.
- [12] Doris Hinestroza and Diego A. Murio, *Identification of transmissivity coefficients by mollification techniques. Part I: one-dimensional Elliptic and Parabolic Problems*, Computers Math. Applic. Vol.25, 59-79, 1993.

6. Referencias

- [13] Eric Ávila Vales, *Ecuaciones de reacción-Difusión*, Universidad Autónoma de Yucatán, Facultad de Matemáticas.
- [14] Friedman, *Partial Differential Equations of Parabolic Type*, Prentice-Hall, 1964.
- [15] Jenifer Peralta, *Identificación del coeficiente la ecuación parabólica unidimensional*, Tesis 2008.
- [16] L. Evans *Partial differential Equations*, American Mathematical society, Volumen 19.
- [17] Engl-Hanke-Neubauer, *Regularization of inverse problems*, 1996.
- [18] Shumakow, N.V., *A method for the experimental study of the process of heating a solid body*, Soviet physics-thechnical physics 2, 771, 1957.
- [19] <http://www.resnet.wm.edu/jxshix/math490/lecture-chap4.pdf>.
- [20] http://en.wikipedia.org/wiki/John_Gordon_Skellam.