



Método Wavelet-Petrov-Galerkin en la solución
numérica de la ecuación KdV

Trabajo de investigación presentado como requisito parcial para optar el
título de Magíster en Matemáticas Aplicadas

Julio César Duarte V

Esper Andrés Fierro Y

Directores

Jairo Villegas G

Jorge Iván Castaño B

Departamento de Ciencias Básicas
Escuela de Ciencias y Humanidades.
Universidad EAFIT
Medellín

Agradecimientos

ESPER ANDRES FIERRO YAGUARA

Agradecer a los directores, Jairo Villegas G y Jorge Iván Castaño B por sus conocimientos y apoyo en todo momento durante la realización de esta Tesis. También brindamos un especial agradecimiento a todos los profesores que tuvimos en la maestría en Matemáticas Aplicadas de la Universidad EAFIT de Medellín los cuales nos ofrecieron valiosos aportes en nuestra formación profesional.

JULIO CESAR DUARTE VIDAL

Agradecer a Dios y a mi familia por esperar y animar mi proyecto de vida, a mis padres por su confianza brindada, a los directores de este trabajo Jairo Villegas G y Jorge Iván Castaño B por sus grandes enseñanzas en el campo investigativo, formación profesional y personal, a la Universidad EAFIT que forma en la investigación científica, a todos nuestros maestros de la maestría por el conocimiento enseñado.

Índice general

Introducción	1
1. Preliminares	5
1.1. Introducción	5
1.2. Transformada de Fourier	7
1.2.1. Serie de Fourier	9
1.3. Distribuciones y espacios de Sobolev	10
1.3.1. Espacios de Sobolev	13
2. Introducción a las wavelets	15
2.1. Introducción	15
2.2. Transformadas wavelets	16
2.2.1. Transformada wavelet continua	16
2.2.2. Transformada wavelet discreta	20
2.3. Análisis Multirresolución	22
2.4. Ecuación de escala	27
2.5. Construcción de la función de escala	29
2.6. Descomposición y reconstrucción	33
2.6.1. Algoritmo de descomposición	34
2.6.2. Algoritmo de reconstrucción	37
2.7. Wavelet biortogonales	38

3. La ecuación de Korteweg-de Vries (KdV)	47
3.1. Introducción	47
3.2. Ondas dispersivas	48
3.3. Derivación de la ecuación KdV	51
3.4. El método de Lax	55
3.5. Método de Adomian para la ecuación KdV	58
4. Método Wavelet-Petrov-Galerkin para la ecuación KdV	63
4.1. Introducción	63
4.2. Método de Petrov-Galerkin para la ecuación KdV	66
4.3. La ecuación KdV linealizada	70
4.4. Convergencia y estabilidad	76
4.4.1. Resultados de convergencia	76
4.4.2. Condiciones de estabilidad	80
4.5. Algoritmos para calcular $a(k)$, $b(l, k)$ y $c(k)$	82
4.6. Implementación del método y resultados numéricos	87
Bibliografía	95

Introducción

Sin lugar a dudas los métodos wavelets¹ permiten desarrollar algoritmos eficientes y novedosos en el estudio del procesamiento de imágenes y señales. La idea de utilizar esta teoría en la solución numérica de ecuaciones en derivadas parciales se da en virtud a que algunas propiedades de las wavelets son importantes en la construcción de algoritmos adaptativos. Un algoritmo de este tipo selecciona un conjunto minimal de aproximaciones en cada paso, de tal manera que la solución calculada sea lo suficientemente próxima a la solución exacta. Si queremos que la solución calculada sea suave en alguna región, sólo unos pocos coeficientes wavelet serán necesarios para obtener una buena aproximación de la solución en dicha región, es decir, solamente los coeficientes de bajas frecuencias cuyo soporte esté en esa región son los utilizados. De otro lado, los coeficientes grandes (en valor absoluto) se localizan cerca de las singularidades y esto nos permite definir criterios de adaptabilidad a través del tiempo de evaluación [15, 23, 53, 64]. Este trabajo se dirige fundamentalmente a encontrar soluciones aproximadas a problemas del tipo hiperbólico o parabólicos, utilizando el método wavelet-Galerkin. El trabajo busca dar respuesta problemas que surgen en diferentes áreas de las ciencias e ingeniería.

El desarrollo de técnicas numéricas para la obtención de soluciones aproximadas de ecuaciones diferenciales se ha incrementado en las últimas

¹Se utilizará este término en lugar de ondícula, palabra también usada por algunos autores.

2 Introducción

décadas. Entre estas técnicas están los métodos de elementos finitos, diferencias finitas y en forma más general el método de los residuos ponderados, en particular el método de Galerkin [11, 12, 17, 29, 38, 42, 57]. Es claro que en cada método se realiza un estudio detallado de la eficiencia y convergencia del método. Recientemente métodos wavelets se están aplicando a la solución numérica de ecuaciones en derivadas parciales [64]. Trabajos pioneros en esta dirección son los de Beylkin [6, 7], Dahmen [22, 23], Jaffard [41], Glowinski [34], Tanaka [62], Vasilyev [65].

El trabajo tiene como motivación inicial, estudiar métodos numéricos aplicados a la solución de problemas de valor inicial o de frontera, en particular, utilizar procesos similares al método de los elementos finitos, pero usando nuevas herramientas como es la teoría wavelets [18, 64]. Para tal fin, se estudiará la ecuación de Korteweg-de Vries (KdV) [27, 28, 35, 44]

$$\begin{aligned}\frac{\partial u}{\partial t} + \alpha u \frac{\partial u}{\partial x} + \beta \frac{\partial^3 u}{\partial x^3} &= 0 \\ u(x, 0) &= u_0(x),\end{aligned}$$

donde α y β son constantes que usualmente se fijan como $\alpha = 6$ y $\beta = 1$. En consecuencia, en el presente trabajo se emplea la solución numérica de la ecuación de Korteweg-de Vries (KdV) para estudiar la interacción de soluciones tipo solitón bajo la perspectiva del análisis wavelets. Entre las principales características de los solitones se encuentra su posibilidad de propagarse como ondas de gran amplitud sin dispersión e interactuar entre ellas de forma tal que luego de la interacción cada onda recupera totalmente sus características previas a la interacción tal como si se hubiera tratado de partículas [14, 28, 52]. La ecuación KdV se resuelve numéricamente empleando el método Wavelet-Petrov-Galerkin; como es clásico en estos problemas, su implementación exige la verificación de la convergencia y la determinación de cotas de error para las soluciones.

Finalmente mencionemos también que las wavelets proporcionan un conjunto de herramientas flexibles para detectar problemas prácticos en ciencia e ingeniería. Entre estas herramientas se tienen la transformada wavelet que está asociada con el Análisis Multirresolución de una señal, es decir, a distintos niveles de resolución se tendrá una base de wavelets. Concretamente, cuanto mayor detalle se pretenda obtener en una señal (mayor resolución), mayor número de funciones por unidad de longitud se tendrán en la base de wavelets, véase por ejemplo, [10, 13, 25, 40, 60, 67, 68]. Además, no existe una transformada wavelet única que resuelva todos

los problemas, a partir de la modelación del proceso y de un análisis a priori del tipo de señal tratada y del objetivo que se pretenda, como por ejemplo, compresión, eliminación del ruido, filtrar u otro. Con esta técnica se busca una familia wavelets, digamos Haar, Daubechies, Coiflets,..., que mejor coincida con las características de la señal a estudiar, véase p.e., [3, 16, 25, 31, 36, 49, 50].

El trabajo comienza con la descripción de la terminología necesaria para abordar los métodos aproximados en la solución de ecuaciones diferenciales parciales, partiendo de los resultados básicos del análisis de Fourier y elementos finitos. El capítulo 2 es una introducción a la teoría wavelet, prestando especial interés a las transformadas wavelets tanto discretas como continuas y terminando con un corto estudio del análisis multirresolución y wavelets biortogonales. En el capítulo 3 se estudia la ecuación de Korteweg-de Vries (KdV) y sus principales propiedades, así como diferentes métodos de solución, esto con el fin que la monografía sea autocontenida. Finalmente en el capítulo 4, se desarrolla el método Wavelet-Petrov-Galerkin aplicado a la ecuación KdV.

CAPÍTULO 1

Preliminares

1.1. Introducción

En este corto capítulo se presentará alguna terminología necesaria para la lectura de esta monografía. En particular, se hará un resumen de resultados básicos de análisis de Fourier omitiendo sus pruebas, las cuales se pueden encontrar en algunos de los siguientes textos [10], [32], [54], [67], [68].

Recuerde que $L_1(\mathbb{R})$ es el espacio de todas las funciones $f : \mathbb{R} \rightarrow \mathbb{C}$, tal que $\int_{\mathbb{R}} |f(t)| dt = \|f\|_{L_1} < \infty$. De igual forma se tiene $L_2(\mathbb{R})$, el espacio de las funciones cuadrado-integrables, cuya norma es

$$\|f\|_{L_2} = \left(\int_{\mathbb{R}} |f(t)|^2 dt \right)^{1/2} < \infty.$$

Este espacio se dota con el producto escalar

$$\langle f, g \rangle_{L_2} = \int_{\mathbb{R}} f(t) \overline{g(t)} dt,$$

donde $\overline{g(t)}$ denota el conjugado complejo de $g(t)$. Con este producto interno el espacio $L_2(\mathbb{R})$ es de Hilbert. Las funciones $f, g \in L_2(\mathbb{R})$ son ortogonales si $\langle f, g \rangle_{L_2} = 0$. En general, $L_p(\mathbb{R})$ ($p \geq 1$), es el espacio de todas las funciones

6 Preliminares

(clases de equivalencia) $f : \mathbb{R} \rightarrow \mathbb{C}$, tal que $\int_{\mathbb{R}} |f(t)|^p dt = \|f\|_{L_p}^p < \infty$, donde

$$\|f\|_{L_p} = \left(\int_{\mathbb{R}} |f(t)|^p dt \right)^{1/p}$$

es la norma de f en $L_p(\mathbb{R})$. Otro espacio que se utilizará es $\ell_2(\mathbb{Z})$, el de las sucesiones (x_j) , $j \in \mathbb{Z}$, tal que $\sum_j |x_j|^2 < \infty$.

Sea $F = \mathbb{C}$ o $\mathbb{R}$, X y Y espacios normados (espacios vectoriales equipados con una norma). Un operador lineal es una función $T : X \rightarrow Y$ tal que $T(au + bv) = aT(u) + bT(v)$, para cada $a, b \in F$ y cada $u, v \in X$. El operador T es continuo en u_0 si para cada $\epsilon > 0$ existe $\delta > 0$ tal que si

$$\|u - u_0\|_X < \delta \quad \text{entonces} \quad \|Tu - Tu_0\|_Y < \epsilon. \quad (1.1.1)$$

Si (1.1.1) se cumple para cada $u_0 \in X$ se dice que T es continuo en X . Si δ no depende del punto u_0 se dice que T es uniformemente continuo en X .

El operador T es acotado si y sólo si existe una constante $c > 0$ tal que $\|Tu\|_Y \leq c\|u\|_X$ para cada $u \in X$.

Si $f, g \in L_1(\mathbb{R})$, entonces la convolución de f y g , denotada $f * g$, se define por

$$(f * g)(t) = \int_{\mathbb{R}} f(t - z)g(z)dz.$$

Un sistema de funciones $\{\phi_j, j \in \mathbb{Z}\}$, $\phi_j \in L_2(\mathbb{R})$, se llama ortonormal si

$$\int_{\mathbb{R}} \phi_j(t) \overline{\phi_k(t)} dt = \delta_{jk},$$

donde δ_{jk} es la delta de Kronecker. Es decir,

$$\delta_{jk} = \begin{cases} 1, & \text{si } j = k; \\ 0, & \text{si } j \neq k. \end{cases}$$

Un sistema ortonormal se llama una base en un subespacio V de $L_2(\mathbb{R})$ si cualquier función $f \in V$ tiene una representación de la forma

$$f(t) = \sum_j c_j \phi_j(t),$$

donde los coeficientes c_j satisfacen $\sum_j |c_j|^2 < \infty$. En lo que sigue se utilizará la notación $\sum_j = \sum_{j=-\infty}^{\infty}$, $\int_{\mathbb{R}} = \int_{-\infty}^{\infty}$, $\|f\|_{L_2} = \|f\|_2$ y $\langle \cdot, \cdot \rangle_2$.

La función característica del conjunto A , χ_A , se define por

$$\chi_A(t) = \begin{cases} 1, & t \in A; \\ 0, & t \notin A. \end{cases}$$

También se utilizará la notación $I\{A\}$ para denotar esta función y la llaman función indicadora.

El soporte de una función $f : A \subseteq \mathbb{R} \rightarrow \mathbb{C}$, denotado $\text{sop}f$, se define por $\text{sop}f = \overline{\{x \in A : f(x) \neq 0\}}$.

1.2. Transformada de Fourier

En esta sección se recordará la definición y algunas propiedades importantes de la transformada de Fourier.

Definición 1.2.1. Sea $f \in L_1(\mathbb{R})$ y $\omega \in \mathbb{R}$. La transformada de Fourier de f en ω se define por

$$\hat{f}(\omega) := \int_{\mathbb{R}} f(t) e^{-i\omega t} dt.$$

Como

$$\int_{\mathbb{R}} |f(t)| |e^{-it\omega}| dt = \int_{\mathbb{R}} |f(t)| dt = \|f\|_{L_1} < \infty$$

se tiene que la transformada de Fourier está bien definida. La aplicación $f \mapsto \hat{f}$ se llama transformación de Fourier y se denota por $\mathcal{F}$ ($\mathcal{F}(f) = \hat{f}$). La función $\hat{f}$ es continua y tiende a cero cuando $|\omega| \rightarrow \infty$ (Lema de Riemann-Lebesgue). Es claro que $\mathcal{F}(af + bg) = a\mathcal{F}(f) + b\mathcal{F}(g)$, para cada $a, b \in \mathbb{R}$.

En general $\hat{f}$ no es una función integrable, por ejemplo, sea

$$f(t) = \begin{cases} 1, & |t| < 1; \\ 0, & |t| > 1. \end{cases}$$

Entonces

$$\begin{aligned} \hat{f}(\omega) &= \int_{-1}^1 e^{-it\omega} dt = \left[\frac{e^{-i\omega} - e^{i\omega}}{-i\omega} \right] \\ &= \frac{2 \operatorname{sen} \omega}{\omega} \notin L_1(\mathbb{R}). \end{aligned}$$

8 Preliminares

Si $\hat{f}(\omega)$ es integrable, entonces existe una versión continua de f y se puede obtener la fórmula de inversión de Fourier

$$f(t) = \mathcal{F}^{-1}(\hat{f}(\omega)) = \frac{1}{2\pi} \int_{\mathbb{R}} \hat{f}(\omega) e^{i\omega t} d\omega. \quad (1.2.1)$$

La siguiente proposición recoge algunas propiedades fundamentales de la transformada de Fourier.

Proposición 1.2.1. Sean $f, g \in L_1(\mathbb{R})$, entonces

1. $\widehat{(\tau_x f)}(\omega) = e^{-i\omega x} \hat{f}(\omega)$, donde $(\tau_a f)(t) = f(t - a)$.
2. $(\tau_x \hat{f})(\omega) = \widehat{(e^{ix(\cdot)} f)}(\omega)$
3. $\widehat{f * g} = \hat{f} \hat{g}$
4. Si $\epsilon > 0$ y $g_\epsilon(t) = g(\epsilon t)$ entonces $\hat{g}_\epsilon(\omega) = \epsilon^{-1} \hat{g}(\omega/\epsilon)$.

Otro resultado útil es el siguiente: Si $f, g \in L_1(\mathbb{R}) \cap L_2(\mathbb{R})$, entonces

$$\|f\|_2^2 = \frac{1}{2\pi} \int_{\mathbb{R}} |\hat{f}(\omega)|^2 d\omega \quad (\text{fórmula de Plancherel}) \quad (1.2.2)$$

$$\langle f, g \rangle_2 = \frac{1}{2\pi} \int_{\mathbb{R}} \hat{f}(\omega) \overline{\hat{g}(\omega)} d\omega \quad (\text{fórmula de Parseval}). \quad (1.2.3)$$

Por extensión, la transformada de Fourier se puede definir para cualquier $f \in L_2(\mathbb{R})$. En virtud a que el espacio $L_1(\mathbb{R}) \cap L_2(\mathbb{R})$ es denso en $L_2(\mathbb{R})$. Luego, por isometría (excepto por el factor $1/2\pi$) se define $\hat{f}$ para cualquier $f \in L_2(\mathbb{R})$, y las fórmulas (1.2.2) y (1.2.3) permanecen válidas para todo $f, g \in L_2(\mathbb{R})$.

En teoría de señales, la cantidad $\|f\|_2$ mide la energía de la señal, mientras que $\|\hat{f}\|_2$ representa el espectro de potencia de f .

Si f es tal que $\int_{\mathbb{R}} |t|^k |f(t)| dt < \infty$, para algún entero $k \geq 1$, entonces

$$\frac{d^k}{d\omega^k} \hat{f}(\omega) = \int_{\mathbb{R}} (-it)^k e^{-i\omega t} f(t) dt. \quad (1.2.4)$$

Recíprocamente, si $\int_{\mathbb{R}} |\omega|^k |\hat{f}(\omega)| d\omega < \infty$, entonces

$$\mathcal{F}(f^{(k)})(\omega) = (i\omega)^k \hat{f}(\omega). \quad (1.2.5)$$

1.2.1. Serie de Fourier

Sea f una función 2π -periódica en $\mathbb{R}$. Se escribirá $f \in L_p(0, 2\pi)$ si

$$f(t)\chi_{[0,2\pi]}(t) \in L_p(0, 2\pi), \quad p \geq 1.$$

Cualquier función f , 2π -periódica en $\mathbb{R}$, tal que $f \in L_2(0, 2\pi)$, se puede representar por una serie de Fourier convergente en $L_2(0, 2\pi)$

$$f(t) = \sum_n c_n e^{int},$$

donde los coeficientes de Fourier son dados por

$$c_n = \frac{1}{2\pi} \int_0^{2\pi} f(t) e^{-int} dt.$$

Se puede verificar que si $f \in L_1(\mathbb{R})$, entonces la serie, fórmula de sumación de Poisson,

$$S(t) = \sum_{k=-\infty}^{\infty} f(t + 2k\pi) = \frac{1}{2\pi} \sum_{n=-\infty}^{\infty} \hat{f}(n) e^{int} \quad (1.2.6)$$

converge casi para todo t y pertenece a $L_1(0, 2\pi)$. Además, los coeficientes de Fourier de $S(t)$ están dados por

$$c_k = \frac{1}{2\pi} \hat{f}(k) = \mathcal{F}^{-1}(f)(-k).$$

En efecto, para ver la expresión (1.2.6), basta probar que

$$\int_0^{2\pi} \sum_k |f(t + 2k\pi)| dt < \infty.$$

Para la segunda parte se calculan los coeficientes de Fourier de $S(t)$, que son los valores de la transformada de Fourier de f en los enteros. Esto es, sea

$$h(t) = \sum_{k=-\infty}^{\infty} f(t + 2k\pi),$$

10 Preliminares

entonces h es 2π -periódica y además, sus coeficientes de Fourier son

$$\begin{aligned}\hat{h}_n &= \frac{1}{2\pi} \int_0^{2\pi} h(t) e^{-int} dt = \frac{1}{2\pi} \int_0^{2\pi} \left[\sum_{k=-\infty}^{\infty} f(t + 2k\pi) \right] e^{-int} dt \\ &= \sum_{k=-\infty}^{\infty} \frac{1}{2\pi} \int_0^{2\pi} f(t + 2k\pi) e^{-int} dt \\ &= \sum_{k=-\infty}^{\infty} \frac{1}{2\pi} \int_{2\pi k}^{2\pi(k+1)} f(z) e^{-in(z-2k\pi)} dz \\ &= \frac{1}{2\pi} \int_{-\infty}^{\infty} f(z) e^{-inz} dz = \frac{1}{2\pi} \hat{f}(n).\end{aligned}$$

Como consecuencia de la fórmula de sumación de Poisson tenemos

$$\sum_{k=-\infty}^{\infty} \hat{f}(\omega + 2k\pi) = \frac{1}{2\pi} \sum_{n=-\infty}^{\infty} f(n) e^{-in\omega} \quad (1.2.7)$$

donde f es una función tal que $\hat{f} \in L_1(\mathbb{R})$, continua, y $\sum_{n=-\infty}^{\infty} f(n)$ converge absolutamente.

1.3. Distribuciones y espacios de Sobolev

En esta sección recogemos algunos resultados básicos sobre distribuciones. La teoría de distribuciones libera al cálculo diferencial de ciertas dificultades que provienen del hecho de que existen funciones no diferenciables. Este hecho extiende el cálculo a una clase de objetos llamados distribuciones o funciones generalizadas, que es mucho mayor que la clase de funciones diferenciables.

Sea Ω un abierto de $\mathbb{R}^n$,

$$\mathcal{D}(\Omega) = C_0^\infty(\Omega) = \{\varphi \in C^\infty(\Omega) : \text{sop } \varphi \text{ es un compacto contenido en } \Omega\},$$

donde $\text{sop } \varphi = \overline{\{x \in \Omega : \varphi(x) \neq 0\}}$. $\mathcal{D}(\Omega)$ denota el espacio vectorial de las funciones de prueba.

Si K es un compacto de Ω entonces

$$\mathcal{D}_K(\Omega) = \{\varphi \in C^\infty(\Omega) : \text{sop } \varphi \subset K\}.$$

Una función típica de $\mathcal{D}(\Omega)$ es

$$\varphi(x) = \begin{cases} 0 & \text{si } |x| \geq 1 \\ c \exp(\frac{1}{|x|^2-1}) & \text{si } |x| < 1, \end{cases}$$

donde $x = (x_1, \dots, x_n) \in \mathbb{R}^n$ y $|x| = \sqrt{x_1^2 + \dots + x_n^2}$, la constante c se escoge de tal forma que $\int_{\mathbb{R}^n} \varphi(x) dx = 1$; $\varphi \in C^\infty(\mathbb{R}^n)$ y su soporte es la bola unitaria en $\mathbb{R}^n$, es decir, $\text{sop } \varphi = B_1(0)$.

Una expresión de la forma $\alpha = (\alpha_1, \dots, \alpha_n)$ con α_i un entero no negativo, para cada $i = 1, \dots, n$, se llama un multi-índice. Para el multi-índice α definimos el orden de α , denotado $|\alpha|$, por

$$|\alpha| = \alpha_1 + \dots + \alpha_n.$$

Si $x = (x_1, \dots, x_n) \in \mathbb{R}^n$ y $\alpha = (\alpha_1, \dots, \alpha_n)$ es un multi-índice, entonces definimos

$$\partial^\alpha u := \frac{\partial^{|\alpha|} u}{\partial x_1^{\alpha_1} \dots \partial x_n^{\alpha_n}}, \quad x^\alpha = x_1^{\alpha_1} x_2^{\alpha_2} \dots x_n^{\alpha_n}.$$

CONVERGENCIA EN $\mathcal{D}(\Omega)$: Sea $(\varphi_j)_{j=1}^\infty$ una sucesión de funciones en $\mathcal{D}(\Omega)$, $\varphi_j \rightarrow \varphi \in \mathcal{D}(\Omega)$, cuando $j \rightarrow \infty$ si:

- a) existe un compacto $K \subset \Omega$ tal que para cada j $\text{sop } \varphi_j \subset K$
- b) $\partial^\alpha \varphi_j \rightarrow \partial^\alpha \varphi$ uniformemente en K , para cada multi-índice α .

La convergencia uniforme en K de la sucesión $(\partial^\alpha \varphi_j)_{j=1}^\infty$ significa que

$$\sup_{x \in K} |(\partial^\alpha \varphi_j - \partial^\alpha \varphi)(x)| \rightarrow 0,$$

cuando $j \rightarrow \infty$. Una aplicación f es continua en $\mathcal{D}(\Omega)$ significa que para cada sucesión $(\varphi_j)_1^\infty$ con límite φ , se tiene $\langle f, \varphi_j \rangle \rightarrow \langle f, \varphi \rangle$, cuando $j \rightarrow \infty$.

Definición 1.3.1. Sea Ω un abierto de $\mathbb{R}^n$. La aplicación $T : \mathcal{D}(\Omega) \rightarrow \mathbb{C}$ es una distribución si

- a) T es lineal.
- b) Para cada compacto $K \subset \Omega$ existe una constante $C_K > 0$ y un entero no negativo m (depende de K) tal que

$$|\langle T, \varphi \rangle| \leq C_K \sum_{|\alpha| \leq m} \sup_{x \in K} |\partial^\alpha \varphi(x)|,$$

para cada $\varphi \in \mathcal{D}_K(\Omega)$ y para cada multi-índice α .

12 Preliminares

En otras palabras, una distribución es un funcional lineal y continuo sobre $\mathcal{D}(\Omega)$. El espacio de todas las distribuciones sobre Ω se denota por $\mathcal{D}'(\Omega)$. Es decir, $\mathcal{D}'(\Omega) = \mathcal{L}(\mathcal{D}(\Omega), \mathbb{C})$ es el dual de $\mathcal{D}(\Omega)$.

Un ejemplo típico de distribución es la delta de Dirac, esto es, considere el funcional lineal $\delta : \mathcal{D}(\mathbb{R}) \rightarrow \mathbb{R}$ definido por $\langle \delta, \varphi \rangle = \varphi(0)$, para cada $\varphi \in \mathcal{D}(\mathbb{R})$. En efecto, si K es un compacto de $\mathbb{R}$ y $\varphi \in \mathcal{D}(\mathbb{R})$ entonces

$$|\langle \delta, \varphi \rangle| = |\varphi(0)| \leq \max_{x \in K} |\varphi(x)|,$$

acá $C_K = 1$ y $m = 0$.

Algunas propiedades importantes de las distribuciones tales como la multiplicación de una función por una distribución y la derivada generalizada se definen a continuación.

Multiplicación de una función $u \in C^\infty(\Omega)$ por una distribución T : Para cada $\varphi \in \mathcal{D}(\Omega)$ se define uT por

$$\langle uT, \varphi \rangle = \langle T, u\varphi \rangle,$$

uT está bien definida ya que si $\varphi \in \mathcal{D}(\Omega)$ y $u \in C^\infty(\Omega)$ entonces $u\varphi \in \mathcal{D}(\Omega)$.

Derivada de una distribución: Si α un multi-índice y $T \in \mathcal{D}'(\Omega)$ se define la derivada de T para cada $\varphi \in \mathcal{D}(\Omega)$ por

$$\langle \partial^\alpha T, \varphi \rangle = (-1)^{|\alpha|} \langle T, \partial^\alpha \varphi \rangle.$$

Con el propósito de extender la transformada de Fourier a las distribuciones, definamos las funciones de decrecimiento rápido.

Definición 1.3.2. Sea $\varphi \in C^\infty(\mathbb{R}^n)$. φ es una función de decrecimiento rápido si para cada α y β multi-índices, existe una constante positiva M tal que

$$|x^\alpha \partial^\beta \varphi(x)| \leq M, \quad \forall x \in \mathbb{R}^n.$$

El conjunto de todas las funciones de decrecimiento rápido forma un espacio vectorial real (o complejo) y lo denotamos por $\mathcal{S}(\mathbb{R}^n)$ y se llama espacio de Schwartz. Los elementos de este espacio se llaman *funciones de prueba de decrecimiento rápido*.

CONVERGENCIA EN $\mathcal{S}(\mathbb{R}^n)$: Una sucesión (φ_j) converge a 0 en $\mathcal{S}(\mathbb{R}^n)$ si y sólo si $x^\alpha \partial^\beta \varphi_j(x) \rightarrow 0$ uniformemente en $\mathbb{R}^n$ cuando $j \rightarrow \infty$.

El dual topológico $\mathcal{S}'(\mathbb{R}^n) := \mathcal{L}(\mathcal{S}(\mathbb{R}^n), \mathbb{C})$ se llama espacio de las distribuciones temperadas.

El funcional lineal $T : \mathcal{S}(\mathbb{R}^n) \rightarrow \mathbb{C}$ es continuo si para cada sucesión (φ_j) tal que $\varphi_j \rightarrow \varphi$ en $\mathcal{S}(\mathbb{R}^n)$ se tiene $\langle T, \varphi_j \rangle \rightarrow \langle T, \varphi \rangle$ para cada $\varphi \in \mathcal{S}(\mathbb{R}^n)$.

1.3.1. Espacios de Sobolev

En este apartado introducimos los espacios de Sobolev y algunas de sus propiedades más importantes. Si queremos estudiar la regularidad de una función de soporte compacto o de una distribución es usual analizar el comportamiento de su transformada de Fourier en el infinito. Una forma alterna de hacer este análisis es midiendo la diferenciabilidad en términos de normas de L_2 . La razón son dos:

- i) L_2 es un espacio de Hilbert.
- ii) La transformada de Fourier, la cual convierte diferenciación en multiplicación por polinomios, es una isometría (isomorfismo que preserva normas).

Definición 1.3.3. Sea $m \geq 0$, $p \geq 1$ y Ω un dominio de $\mathbb{R}^n$ ($n \geq 2$). El espacio de Sobolev $W^{m,p}(\Omega)$ se define como

$$W^{m,p}(\Omega) = \{u \in L_p(\Omega) : \partial^\alpha u \in L_p(\Omega), \quad \forall \alpha, \quad |\alpha| \leq m\}.$$

Es un espacio vectorial normado equipado con la norma de Sobolev

$$\|u\|_{W^{m,p}(\Omega)} = \left(\int_{\Omega} \sum_{|\alpha| \leq m} |\partial^\alpha u|^p dx \right)^{1/p} = \left(\sum_{|\alpha| \leq m} \|\partial^\alpha u\|_{L_p(\Omega)}^p \right)^{1/p}.$$

Propiedades

1. $W^{m,p}(\Omega) = \widetilde{C^m(\Omega)}$ (completado) en la norma de Sobolev $\|\cdot\|_{W^{m,p}(\Omega)}$.
2. $(W^{m,p}(\Omega), \|\cdot\|_{W^{m,p}(\Omega)})$ es un espacio de Banach.
3. $\overline{C_0^m(\Omega)}^{W^{m,p}(\Omega)} = W_0^{m,p}(\Omega) = \{u \in W^{m,p}(\Omega) : \partial^\alpha u|_{\partial\Omega} = 0, \quad |\alpha| \leq m-1\}$, $W_0^{m,p}(\Omega)$ es un subespacio de $W^{m,p}(\Omega)$ (acá $\partial^\alpha u|_{\partial\Omega}$ denota la extensión de $\partial^\alpha u$ a la frontera de Ω).
4. $(W_0^{m,p}(\Omega))' := W^{-m,q}(\Omega)$, donde $\frac{1}{p} + \frac{1}{q} = 1$, se llaman espacios de Sobolev negativos (si $p = 1$, $(W^{m,1}(\Omega))' = W^{-m,\infty}(\Omega)$).
5. Si $u \in W_0^{m,p}(\Omega)$ y $v \in W^{-m,q}(\Omega)$ es integrable, entonces

$$\left| \int_{\Omega} u v dx \right| \leq \|u\|_{W^{m,p}(\Omega)} \|v\|_{W^{-m,q}(\Omega)}, \quad \frac{1}{p} + \frac{1}{q} = 1.$$

14 Preliminares

6. Si $m > k$ entonces

$$W^{m,p}(\Omega) \subset W^{k,p}(\Omega) \quad \text{y} \quad \|\cdot\|_{W^{k,p}(\Omega)} \leq \|\cdot\|_{W^{m,p}(\Omega)},$$

en particular, $W^{0,p}(\Omega) = L_2(\Omega)$ y $\|\cdot\|_{W^{0,p}(\Omega)} = \|\cdot\|_p$.

Si $p = 2$, en lugar de $W^{m,2}(\Omega)$ escribiremos $H^m(\Omega)$, o sea,

$$H^m(\Omega) = \{u \in L_2(\Omega) : \partial^\alpha u \in L_2(\Omega), \quad |\alpha| \leq m\},$$

$H^m(\Omega)$ es un espacio de Hilbert con el producto interno definido por

$$\langle u, v \rangle_{H^m(\Omega)} = \int_{\Omega} \sum_{|\alpha| \leq m} \partial^\alpha u \partial^\alpha v dx,$$

$u, v \in H^m(\Omega)$. La norma en $H^m(\Omega)$ es

$$\|u\|_{H^m(\Omega)} = \sqrt{\langle u, u \rangle_{H^m(\Omega)}} = \left(\sum_{|\alpha| \leq m} \|\partial^\alpha u\|_{L_2(\Omega)}^2 \right)^{1/2}.$$

Para $s \in \mathbb{R}$, el espacio de Sobolev $H^s(\mathbb{R})$ se define por

$$H^s(\mathbb{R}) := \left\{ f \in L_2(\mathbb{R}) : \int_{\mathbb{R}} (1 + |x|^2)^s |\hat{f}(x)|^2 dx < \infty \right\},$$

donde $\hat{f}$ es la transformada de Fourier de f . $H^s(\mathbb{R})$ equipado con el producto interno

$$\langle f, g \rangle_{H^s} := \int_{\mathbb{R}} (1 + |x|^2)^s \hat{f}(x) \overline{\hat{g}(x)} dx$$

es un espacio de Hilbert. También se le asocia la norma

$$\|f\|_{H^s} := \left(\int_{\mathbb{R}} (1 + |x|^2)^s |\hat{f}(x)|^2 dx \right)^{1/2}.$$

La notación O

La expresión $f(x) = O(g(x))$ significa que existe una constante positiva M tal que $|f(x)| \leq M|g(x)|$, siempre que $x \rightarrow x_0$. Es decir, si $g(x) \neq 0$ entonces $\left| \frac{f(x)}{g(x)} \right| \rightarrow M$, cuando $x \rightarrow x_0$.

CAPÍTULO 2

Introducción a las wavelets

2.1. Introducción

El origen de la descomposición de una señal en wavelets está en la necesidad de conocer las características y particularidades de la señal en diferentes instantes de tiempo. La principal virtud de las wavelets es que permite modelar procesos que dependen fuertemente del tiempo y para los cuales su comportamiento no tiene porqué ser suave [3], [16]. Una de las ventajas de las wavelets frente a los métodos clásicos, como la transformada de Fourier, es que en el segundo caso se maneja una base de funciones bien localizada en frecuencia pero no en tiempo, esto es, el análisis en frecuencia obtenido del análisis de Fourier es insensible a perturbaciones que supongan variaciones instantáneas y puntuales de la señal como picos debidos a conmutaciones o variaciones muy lentas como tendencias. En otras palabras, si f es una señal (f es una función definida en todo $\mathbb{R}$ y tiene energía finita $\int_{-\infty}^{\infty} |f(t)|^2 dt$). La transformada de Fourier $\hat{f}(\omega)$ proporciona la información global de la señal en el tiempo localizada en frecuencia. Sin embargo, $\hat{f}(\omega)$ no particulariza la información para intervalos de tiempo específicos, ya que

$$\hat{f}(\omega) = \int_{-\infty}^{\infty} f(t)e^{-i\omega t} dt$$

y la integración es sobre todo tiempo (ver [32]). Así, la imagen obtenida no contiene información sobre tiempos específicos, sino que sólo permite calcular el espectro de amplitud total $|\hat{f}(\omega)|$, mientras que la mayoría de las wavelets interesantes presentan una buena localización en tiempo y en frecuencia, disponiendo incluso de bases de wavelets con soporte compacto.

En este capítulo se presenta una introducción a la teoría wavelets, en particular se estudiará la transformada wavelet y el análisis multirresolución en $L_2(\mathbb{R})$. Con este concepto se ilustra como construir otras bases wavelets, y además, permite analizar funciones (señales) en $L_2(\mathbb{R})$ en varias escalas (niveles de resolución) [13], [16], [25]. Para ello, se utiliza versiones escaladas de un conjunto ortonormal en $L_2(\mathbb{R})$. Para tal descomposición de una función $f \in L_2(\mathbb{R})$, sólo se necesitan los coeficientes de la expansión de f en dicho conjunto ortonormal.

2.2. Transformadas wavelets

El análisis wavelets es un método de descomposición de una función o señal usando funciones especiales, las wavelets. La descomposición es similar a la de la transformada de Fourier, donde una señal $f(t)$ se descompone en una suma infinita de armónicos $e^{i\omega t}$ de frecuencias $\omega \in \mathbb{R}$, cuyas amplitudes son los valores de la transformada de Fourier de f , $\hat{f}(\omega)$:

$$f(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\omega) e^{i\omega t} d\omega, \quad \text{donde} \quad \hat{f}(\omega) = \int_{-\infty}^{\infty} f(t) e^{-i\omega t} dt.$$

El análisis de Fourier tiene el defecto de la no localidad: el comportamiento de una función en un conjunto abierto, no importa cuán pequeño, influye en el comportamiento global de la transformada de Fourier. No se captan los aspectos locales de la señal tales como cambios bruscos, saltos o picos, que se han de determinar a partir de su reconstrucción.

2.2.1. Transformada wavelet continua

La teoría wavelets se basa en la representación de una función en términos de una familia biparamétrica de dilataciones y traslaciones de una función fija ψ , la wavelet madre que, en general, no es senoidal. Por ejemplo,

$$f(t) = \int_{\mathbb{R}^2} \frac{1}{\sqrt{|a|}} \psi\left(\frac{t-b}{a}\right) \mathcal{W}_\psi f(a, b) da db$$

en donde $\mathcal{W}_\psi f$ es una transformada de f definida adecuadamente. También se tiene de modo alterno un desarrollo en serie

$$f(t) = \sum_{j,k} c_{j,k} 2^{j/2} \psi(2^j t - k)$$

en donde se suma sobre las dilataciones en progresión geométrica. Para conservar la norma en $L_2(\mathbb{R})$ de la wavelet madre ψ , se insertan los factores $\frac{1}{\sqrt{|a|}}$ y $2^{j/2}$, respectivamente.

Definición 2.2.1. Una wavelet ψ es una función cuadrado integrable tal que la siguiente condición de admisibilidad se tiene

$$C_\psi := \int_{\mathbb{R}} \frac{|\hat{\psi}(\omega)|^2}{|\omega|} d\omega < \infty, \quad (2.2.1)$$

donde $\hat{\psi}(\omega)$ es la transformada de Fourier de ψ .

Observación 2.2.1. Si además $\psi \in L_1(\mathbb{R})$, entonces la condición (2.2.1) implica que $\int_{\mathbb{R}} \psi(t) dt = 0$. En efecto, por el Lema de Riemann-Lebesgue (ver p.e., [54]), $\lim_{\omega \rightarrow \infty} \hat{\psi}(\omega) = 0$ y la transformada de Fourier es continua, lo cual implica que $0 = \hat{\psi}(0) = \int_{\mathbb{R}} \psi(t) dt$.

Sea $\psi \in L_2(\mathbb{R})$. La función dilatada y trasladada se define por

$$\psi_{a,b}(t) := \frac{1}{\sqrt{|a|}} \psi\left(\frac{t-b}{a}\right), \quad a, b \in \mathbb{R}, \quad a \neq 0.$$

Esta función se obtiene a partir de ψ , primero por dilatación en el factor a y, luego, por traslación en b . Es claro que $\|\psi_{a,b}\|_2 = \|\psi\|_2$.

Definición 2.2.2. Para $f, \psi \in L_2(\mathbb{R})$, la expresión

$$\mathcal{W}_\psi f(a, b) := \int_{\mathbb{R}} f(t) \overline{\psi_{a,b}(t)} dt \quad (2.2.2)$$

se llama la transformada wavelet de f .

Por la desigualdad de Cauchy, se ve que $\mathcal{W}_\psi f$ es una función acotada con $|\mathcal{W}_\psi f(a, b)| \leq \|f\|_2 \|\psi\|_2$. Note también que

$$\mathcal{W}_\psi f(a, b) = \langle f, \psi_{a,b} \rangle_{L_2(\mathbb{R})} = \langle f, \psi_{a,b} \rangle.$$

18 Introducción a las wavelets

La transformada wavelet $\mathcal{W}_\psi f$ de f puede ser descrita en términos del producto de convolución. La convolución de dos funciones $f, g \in L_2(\mathbb{R})$ es dada por

$$(f * g)(t) = \int_{\mathbb{R}} f(t - z)g(z)dz.$$

Observe que esta fórmula está definida para al menos todo $t \in \mathbb{R}$, pero $f * g$ no necesariamente está en $L_2(\mathbb{R})$. Usando la notación $\tilde{\psi}(t) = \overline{\psi(-t)}$, se tiene $\mathcal{W}_\psi f(a, b) = (f * \tilde{\psi}_{a,0})(b)$. Note también que $\hat{\psi}_{a,b}(\omega) = \sqrt{|a|}\tilde{\hat{\psi}}(a\omega)e^{-i\omega b}$. Estos hechos se aplicarán en la prueba de la siguiente proposición, la cual establece la fórmula de Plancherel para la transformada wavelet.

Proposición 2.2.1. *Sea $\psi \in L_2(\mathbb{R})$ y satisface la condición (2.2.1). Entonces para cualquier $f \in L_2(\mathbb{R})$, las siguientes relaciones se tienen*

1. *Isometría*

$$\int_{\mathbb{R}} |f(t)|^2 dt = \frac{1}{C_\psi} \int_{\mathbb{R}^2} |\mathcal{W}_\psi f(a, b)|^2 db \frac{da}{a^2}$$

2. *Fórmula de inversión*

$$f(t) = \frac{1}{C_\psi} \int_{\mathbb{R}^2} \mathcal{W}_\psi f(a, b) \psi_{a,b}(t) db \frac{da}{a^2}$$

Demostración. 1. Es fácil verificar que $(f * \tilde{\psi}_{a,0})(b) = \sqrt{|a|}\mathcal{F}^{-1}\{\hat{f}(\omega)\tilde{\hat{\psi}}(a\omega)\}$. En consecuencia,

$$\begin{aligned} \int_{\mathbb{R}^2} |\mathcal{W}_\psi f(a, b)|^2 db \frac{da}{a^2} &= \int_{\mathbb{R}} \int_{\mathbb{R}} |(f * \tilde{\psi}_{a,0})(b)|^2 db \frac{da}{a^2} \\ &= \int_{\mathbb{R}} \int_{\mathbb{R}} |a| |\mathcal{F}^{-1}(\hat{f}(\cdot)\tilde{\hat{\psi}}(a\cdot))(\omega)|^2 d\omega \frac{da}{a^2} \\ &= \int_{\mathbb{R}} \int_{\mathbb{R}} |\hat{f}(\omega)|^2 |\hat{\psi}(a\omega)|^2 d\omega \frac{da}{|a|} \\ &= \int_{\mathbb{R}} |\hat{f}(\omega)|^2 \left[\int_{\mathbb{R}} |\hat{\psi}(a\omega)|^2 \frac{da}{|a|} \right] d\omega \\ &= C_\psi \int_{\mathbb{R}} |\hat{f}(\omega)|^2 d\omega = C_\psi \|f\|_2^2. \end{aligned}$$

Observe que se utilizó el teorema de Fubini y la fórmula de Plancherel para la transformada de Fourier.

2. Para simplificar los cálculos en la fórmula de inversión, suponga que $f, \hat{f} \in L_1(\mathbb{R})$.

$$\begin{aligned} \int_{\mathbb{R}} \mathcal{W}_\psi f(a, b) \psi_{a,b}(t) db &= \sqrt{|a|} \int_{\mathbb{R}} \mathcal{F}^{-1} \left(\hat{f}(\cdot) \tilde{\psi}(a \cdot) \right) (\omega) \psi_{a,b}(t) d\omega \\ &= \sqrt{|a|} \int_{\mathbb{R}} \hat{f}(\omega) \tilde{\psi}(a \omega) \mathcal{F}^{-1}(g)(\omega) d\omega, \end{aligned}$$

donde $g(b) := \psi_{a,b}(t)$. Ahora, la transformada inversa de Fourier de g es

$$\begin{aligned} \mathcal{F}^{-1}(g)(\omega) &= \frac{1}{2\pi} \int_{\mathbb{R}} g(b) e^{i\omega b} db \\ &= \frac{1}{2\pi} \sqrt{|a|} \int_{\mathbb{R}} \psi(z) e^{-ia\omega z} e^{i\omega t} dz \\ &= \frac{1}{2\pi} \sqrt{|a|} \hat{\psi}(a\omega) e^{i\omega t}. \end{aligned}$$

Sustituyendo e integrando respecto a $a^{-2} da$ se obtiene

$$\begin{aligned} \int_{\mathbb{R}^2} \mathcal{W}_\psi f(a, b) \psi_{a,b}(t) db \frac{da}{a^2} &= \frac{1}{2\pi} \int_{\mathbb{R}} |a| \left[\int_{\mathbb{R}} \hat{f}(\omega) |\hat{\psi}(a\omega)|^2 e^{i\omega t} d\omega \right] \frac{da}{a^2} \\ &= \frac{1}{2\pi} \int_{\mathbb{R}} \hat{f}(\omega) \left[\int_{\mathbb{R}} |\hat{\psi}(a\omega)|^2 \frac{da}{|a|} \right] e^{i\omega t} d\omega \\ &= C_\psi \frac{1}{2\pi} \int_{\mathbb{R}} \hat{f}(\omega) e^{i\omega t} d\omega \\ &= C_\psi f(t). \end{aligned}$$

□

Otro resultado de interés que se presentará en la siguiente proposición, es la fórmula de Parseval para la transformada wavelet.

Proposición 2.2.2. *Sea $\psi \in L_2(\mathbb{R})$ y satisface la condición (2.2.1). Entonces para cualquier $f, g \in L_2(\mathbb{R})$, se tienen*

$$\langle f, g \rangle_{L_2(\mathbb{R})} = \frac{1}{C_\psi} \int_{\mathbb{R}^2} \mathcal{W}_\psi f(a, b) \overline{\mathcal{W}_\psi g(a, b)} \frac{dad b}{a^2}$$

20 Introducción a las wavelets

Demostración. Como $(f * \tilde{\psi}_{a,0})(b) = \sqrt{|a|} \mathcal{F}^{-1}\{\hat{f}(\omega) \tilde{\hat{\psi}}(a\omega)\}$ o de manera equivalente, $\mathcal{F}(f * \tilde{\psi}_{a,0})(\omega) = \sqrt{|a|} \hat{f}(\omega) \tilde{\hat{\psi}}(a\omega)$, entonces

$$\int_{\mathbb{R}} \mathcal{W}_{\psi} f(a, b) \overline{\mathcal{W}_{\psi} g(a, b)} db = |a| \int_{\mathbb{R}} \hat{f}(\omega) \tilde{\hat{g}}(\omega) |\hat{\psi}(a\omega)|^2 d\omega,$$

ahora, integrando respecto a $a^{-2} da$ se sigue

$$\begin{aligned} \int_{\mathbb{R}^2} \mathcal{W}_{\psi} f(a, b) \overline{\mathcal{W}_{\psi} g(a, b)} db \frac{da}{a^2} &= \int_{\mathbb{R}} |a| \left[\int_{\mathbb{R}} \hat{f}(\omega) \tilde{\hat{g}}(\omega) |\hat{\psi}(a\omega)|^2 d\omega \right] \frac{da}{a^2} \\ &= \int_{\mathbb{R}} \hat{f}(\omega) \tilde{\hat{g}}(\omega) \left[\int_{\mathbb{R}} |\hat{\psi}(a\omega)|^2 \frac{da}{|a|} \right] d\omega \\ &= C_{\psi} \int_{\mathbb{R}} \hat{f}(\omega) \tilde{\hat{g}}(\omega) d\omega \\ &= C_{\psi} \langle \hat{f}, \hat{g} \rangle_{L_2(\mathbb{R})} = C_{\psi} \langle f, g \rangle_{L_2(\mathbb{R})}. \end{aligned}$$

Note que se aplicó el teorema de Fubini, y en el último renglón de la expresión anterior, la fórmula de Parseval para la transformada de Fourier. $\square$

En la siguiente proposición se listan algunas propiedades.

Proposición 2.2.3. Sean ψ y φ wavelets y $f, g \in L_2(\mathbb{R})$. Entonces

1. $\mathcal{W}_{\psi}(\alpha f + \beta g)(a, b) = \alpha \mathcal{W}_{\psi} f(a, b) + \beta \mathcal{W}_{\psi} g(a, b)$, $\alpha, \beta \in \mathbb{R}$.
2. $\mathcal{W}_{\alpha\psi + \beta\varphi} f(a, b) = \bar{\alpha} \mathcal{W}_{\psi} f(a, b) + \bar{\beta} \mathcal{W}_{\varphi} f(a, b)$, $\alpha, \beta \in \mathbb{R}$.
3. $\mathcal{W}_{\psi}(\tau_c f)(a, b) = \mathcal{W}_{\psi} f(a, b - c)$, donde τ_c es el operador traslación definido por $\tau_c f(t) = f(t - c)$.
4. $\mathcal{W}_{\psi}(D_c f)(a, b) = \sqrt{c} \mathcal{W}_{\psi} f(ca, cb)$, donde D_c es el operador dilatación definido por $D_c f(t) = \sqrt{c} f(ct)$.

2.2.2. Transformada wavelet discreta

La transformada wavelet continua introduce cierta redundancia, pues la señal original se puede reconstruir completamente calculando $\mathcal{W}_{\psi} f(a, \cdot)$ para una cantidad numerable de escalas, por ejemplo, potencias enteras de 2. Esto es, si se elige la escala $a = 2^{-j}$ para cada $j \in \mathbb{Z}$, y también se discretiza en

el dominio del tiempo en los puntos $b = 2^{-j}k$, $k \in \mathbb{Z}$, la familia de wavelets será ahora dada por

$$\psi_{2^{-j}, 2^{-j}k}(t) = \frac{1}{\sqrt{2^{-j}}} \psi\left(\frac{t - 2^{-j}k}{2^{-j}}\right) = 2^{j/2} \psi(2^j t - k), \quad \forall j, k \in \mathbb{Z}.$$

Se utilizará la notación ψ_{jk} para denotar la wavelet ψ comprimida 2^j y trasladada el entero k , es decir, $\psi_{jk}(t) = 2^{j/2} \psi(2^j t - k)$.

Con la elección de $a = 2^{-j}$ y $b = 2^{-j}k$, observe que el muestreo en el tiempo se ajusta proporcionalmente a la escala, es decir, a mayor escala se toma puntos más distantes, ya que se busca información global, mientras que a menor escala se buscan detalles de la señal, por tal motivo se muestrea en puntos menos distantes entre si. Para otras elecciones de a y b se puede consultar [16].

Definición 2.2.3. Una función $\psi \in L_2(\mathbb{R})$ es una wavelet si la familia de funciones ψ_{jk} definidas por

$$\psi_{jk}(t) = 2^{j/2} \psi(2^j t - k), \quad \forall j, k \in \mathbb{Z}, \quad (2.2.3)$$

es una base ortonormal en el espacio $L_2(\mathbb{R})$.

Una condición suficiente para la reconstrucción de una señal f es que la familia de dilatadas y trasladadas ψ_{jk} forme una base ortonormal en el espacio $L_2(\mathbb{R})$, ver [25] y [36] para más detalles. Si esto se tiene, cualquier función $f \in L_2(\mathbb{R})$ se puede escribir como

$$f(t) = \sum_{j,k} c_{j,k} \psi_{jk}(t) \quad (2.2.4)$$

o teniendo en cuenta (2.2.3) como

$$f(t) = \sum_{j,k} c_{j,k} 2^{j/2} \psi(2^j t - k),$$

donde $c_{j,k} = \langle f, \psi_{2^{-j}, 2^{-j}k} \rangle = \mathcal{W}_\psi f(2^{-j}, 2^{-j}k)$.

Definición 2.2.4. Para cada $f \in L_2(\mathbb{R})$ el conjunto bidimensional de coeficientes

$$c_{j,k} = \langle f, \psi_{jk} \rangle = \int_{\mathbb{R}} 2^{j/2} f(t) \overline{\psi(2^j t - k)} dt$$

se llama la transformada wavelet discreta de f .

22 Introducción a las wavelets

En consecuencia, la expresión (2.2.4) se puede escribir en forma alterna como

$$f(t) = \sum_{j,k} \langle f(t), \psi_{jk}(t) \rangle \psi_{jk}(t). \quad (2.2.5)$$

La serie (2.2.5) se llama representación wavelet de f .

Observación 2.2.2. $\psi_{jk}(t)$ es muy apropiada para representar detalles más finos de la señal como oscilaciones rápidas. Los coeficientes wavelet $c_{j,k}$ miden la cantidad de fluctuaciones sobre el punto $t = 2^{-j}k$ con una frecuencia determinada por el índice de dilatación j .

Es interesante notar que $c_{j,k} = \mathcal{W}_\psi f(2^{-j}, 2^{-j}k)$ es la transformada wavelet de f en el punto $(2^{-j}, 2^{-j}k)$. Estos coeficientes analizan la señal mediante la wavelet madre ψ .

2.3. Análisis Multirresolución

El sistema de Haar no es muy apropiado para aproximar funciones suaves. De hecho, cualquier aproximación de Haar es una función discontinua [25], [36]. Se puede probar que si f es una función muy suave, los coeficientes de Haar decrecerán muy lentamente. Por tanto se pretende construir wavelets que tengan mejor propiedades de aproximación, y una forma de hacerlo es a través del análisis multirresolución (AMR) [50], [49], [48], [47], [25].

Sea $\varphi \in L_2(\mathbb{R})$, la familia de trasladadas de φ ,

$$\{\varphi_{0k}, k \in \mathbb{Z}\} = \{\varphi_{0k}(\cdot - k), k \in \mathbb{Z}\}$$

es un sistema ortonormal (con el producto interno de $L_2(\mathbb{R})$). Acá y en lo que sigue

$$\varphi_{jk}(t) = 2^{j/2} \varphi(2^j t - k) = D_{2^j} \tau_k \varphi(t), \quad j \in \mathbb{Z}, \quad k \in \mathbb{Z},$$

recuerde que $D_a f(t) = a^{1/2} f(at)$ y $\tau_a f(t) = f(t - a)$ son los operadores dilatación y traslación, respectivamente.

Se definen los espacios vectoriales

$$\begin{aligned} V_0 &= \left\{ f(t) = \sum_k c_k \varphi(t - k) : \sum_k |c_k|^2 < \infty \right\}, \\ V_1 &= \left\{ h(t) = f(2t) : f \in V_0 \right\}, \\ &\vdots \\ V_j &= \left\{ h(t) = f(2^j t) : f \in V_0 \right\}, j \in \mathbb{Z} \\ &= \text{gen}\{\varphi_{jk}(t) = 2^{j/2} \varphi(2^j t - k) : k \in \mathbb{Z}\}. \end{aligned}$$

Note que φ genera la sucesión de espacios $\{V_j, j \in \mathbb{Z}\}$. Suponga que la función φ se escoge de tal forma que los espacios estén encajados $V_j \subset V_{j+1}$, $j \in \mathbb{Z}$, y $\bigcup_{j \geq 0} V_j$ es denso en $L_2(\mathbb{R})$, estos dos hechos fundamentales hacen parte de la definición de análisis multirresolución.

Definición 2.3.1. *Un análisis multirresolución en $L_2(\mathbb{R})$ es una sucesión creciente de subespacios cerrados V_j , $j \in \mathbb{Z}$, en $L_2(\mathbb{R})$,*

$$\cdots \subset V_{-2} \subset V_{-1} \subset V_0 \subset V_1 \subset V_2 \subset \cdots$$

tales que

1. $\bigcup_{j \in \mathbb{Z}} V_j$ es denso en $L_2(\mathbb{R})$,
2. $\bigcap_{j \in \mathbb{Z}} V_j = \{0\}$,
3. $f(t) \in V_j \Leftrightarrow f(2t) \in V_{j+1}$, $j \in \mathbb{Z}$,
4. $f(t) \in V_0 \Leftrightarrow f(t - k) \in V_0$, $j \in \mathbb{Z}$,
5. Existe una función $\varphi \in L_2(\mathbb{R})$ tal que el conjunto de funciones $\{\varphi(t - k)\}_{k \in \mathbb{Z}}$ es una base ortonormal para V_0 .

La función φ se llama función de escala. En el espacio V_{j+1} las funciones (señales) se describen con más detalle que en el espacio V_j , la resolución es mejor en el espacio “más grande”. Esto es, las funciones en V_{j+1} que no están en V_j realzan la resolución [16]. Es usual reunir estos “sintonizadores finos” en un nuevo subespacio $W_j = V_{j+1} \setminus V_j$. Sin embargo, la elección de estos subespacios no es única. Pero se puede escoger a W_j como el complemento ortogonal de V_j en V_{j+1} . Es decir,

$$W_j = V_{j+1} \cap V_j^\perp, \quad j \in \mathbb{Z},$$

24 Introducción a las wavelets

o de manera equivalente

$$V_{j+1} = V_j \oplus W_j, \quad j \in \mathbb{Z}. \quad (2.3.1)$$

Informalmente, esto quiere decir que si se tiene una función (señal) f a resolución 2^{j+1} y se proyecta a resolución inferior 2^j entonces

$$f = P_j f + \sum_{k \in \mathbb{Z}} \langle f, \psi_{jk} \rangle \psi_{jk},$$

acá P_j representa la proyección ortogonal en el espacio V_j donde se recoge la versión “suavizada” de f y la diferencia $f - P_j f$ representa el “detalle” de f , que está en W_j y se expresa como $\sum_{k \in \mathbb{Z}} \langle f, \psi_{jk} \rangle \psi_{jk}$. Recuerde que

$$P_j f = \sum_{k \in \mathbb{Z}} \langle f, \varphi_{jk} \rangle \varphi_{jk}, \quad j \in \mathbb{Z}.$$

En otras palabras, W_j contiene los detalles en V_{j+1} que no se representan en V_j , y cada función (señal) en W_j es ortogonal a toda función en V_j (ver p.e., [10]).

El conjunto de funciones linealmente independientes φ_{jk} que generan a V_j son las funciones de escala, mientras que el conjunto de funciones linealmente independientes ψ_{jk} que generan a W_j son las wavelets.

Por definición, el subespacio W_j es cerrado. Note también que si $f \in V_0$, entonces por 5 de la definición anterior se tiene

$$f(t) = \sum_k \langle f, T_k \varphi \rangle T_k \varphi(t).$$

Además, por la ortogonalidad de $\{T_k \varphi(t)\}_{k \in \mathbb{Z}}$,

$$\sum_k |\langle f, T_k \varphi \rangle|^2 = \|f\|_2^2.$$

Observe que al aplicar la descomposición (2.3.1) en cada V_j se obtiene

$$\begin{aligned} V_N &= V_{N-1} \oplus W_{N-1} = V_{N-2} \oplus W_{N-2} \oplus W_{N-1} \\ &= \cdots = V_{-N} \oplus \left(\bigoplus_{j=-N}^{N-1} W_j \right), \end{aligned}$$

y cuando $N \rightarrow \infty$ se tiene

$$\overline{\bigcup_{j \in \mathbb{Z}} V_j} = \bigoplus_{j \in \mathbb{Z}} W_j \oplus \bigcap_{j \in \mathbb{Z}} V_j.$$

Usando las condiciones 1 y 2 de la definición de AMR se obtiene

$$\bigoplus_{j \in \mathbb{Z}} W_j = L_2(\mathbb{R}).$$

Por definición, también los subespacios W_j satisfacen las condiciones 3 y 4 de la definición de AMR o de manera directa como se prueba en el siguiente lema.

Lema 2.3.1. *Sea $(V_j)_{j \in \mathbb{Z}}$ un AMR y $W_j = V_{j+1} \cap V_j^\perp$. Entonces*

- i) $f \in W_j \Leftrightarrow \tau_k f \in W_j$, para cada $k, j \in \mathbb{Z}$.
- ii) $f \in W_j \Leftrightarrow Df \in W_{j+1}$, para cada $j \in \mathbb{Z}$.

Demostración. Sea $f \in W_j$, esto significa que $f \in V_{j+1}$ y $\langle f, D_{2^j} \tau_k \varphi \rangle = 0$ para cada $k \in \mathbb{Z}$. Por la condición 3 y 4 de AMR, la primera relación es equivalente a $\tau_k f \in V_{j+1}$ y $Df \in V_{j+2}$. Además de la relación $\tau_k D_2 = D_2 \tau_k$, se sigue inmediatamente, que la segunda relación es equivalente a

$$\langle \tau_k f, \tau_k D_{2^j} \tau_k \varphi \rangle = \langle \tau_k f, D_{2^j} \tau_{k+2^j} \varphi \rangle = 0, \quad \forall k \in \mathbb{Z}.$$

Por tanto $\tau_k f \in V_{j+1} \cap V_{j+1}^\perp$, y así, $\tau_k f \in W_j$.

La segunda relación también es equivalente con

$$\langle Df, D_{2^{j+1}} \tau_k \varphi \rangle, \quad \forall k \in \mathbb{Z},$$

de lo cual se sigue que $Df \in V_{j+1}^\perp$ que junto con $Df \in V_{j+2}$ se obtiene $Df \in W_{j+1}$. $\square$

La siguiente proposición justifica los comentarios hechos arriba y es útil en futuros resultados.

Proposición 2.3.1. *Sea $(V_j)_{j \in \mathbb{Z}}$ un análisis multirresolución con función de escala φ . Entonces para cada $j \in \mathbb{Z}$, el conjunto de funciones*

$$\{\varphi_{jk}(t) = 2^{j/2} \varphi(2^j t - k), k \in \mathbb{Z}\}$$

es una base ortonormal para V_j .

Demostración. Para probar que $\{\varphi_{jk}(t), k \in \mathbb{Z}\}$ genera a V_j , se debe ver que toda $f(t) \in V_j$ se puede escribir como combinación lineal de funciones de $\{\varphi(2^j t - k), k \in \mathbb{Z}\}$. La propiedad 3 de la definición de AMR, implica que la función $f(2^{-j}t)$ pertenece a V_0 y por tanto $f(2^{-j}t)$ es combinación lineal de $\{\varphi(t - k), k \in \mathbb{Z}\}$. Haciendo la transformación $t \mapsto 2^j t$, se tiene que $f(t)$ es combinación lineal de $\{\varphi(2^j t - k), k \in \mathbb{Z}\}$.

Resta probar que $\{\varphi_{jk}(t), k \in \mathbb{Z}\}$ es ortonormal. Para ello se debe ver que

$$\langle \varphi_{jk}, \varphi_{jm} \rangle = \delta_{jk} = \begin{cases} 0, & \text{si } j \neq k; \\ 1, & \text{si } j = k \end{cases}$$

o

$$2^j \int_{-\infty}^{\infty} \varphi(2^j t - k) \overline{\varphi(2^j t - m)} dt = \delta_{km}.$$

Para establecer esta igualdad, basta hacer el cambio de variable $z = 2^j t$, para obtener

$$2^j \int_{-\infty}^{\infty} \varphi(2^j t - k) \overline{\varphi(2^j t - m)} dt = \int_{-\infty}^{\infty} \varphi(z - k) \overline{\varphi(z - m)} dz = \delta_{km},$$

en virtud de la propiedad 5 de la definición de AMR. □

El siguiente lema contiene dos resultados utilizados en la existencia de los sistemas AMR, bajo hipótesis apropiadas. Por motivos de completitud se hará la prueba del primer apartado, la del segundo se puede encontrar en [67]. Recuerde que P_j es la proyección ortogonal sobre el espacio V_j .

Lema 2.3.2. *Para cualquier $f \in L_2(\mathbb{R})$,*

$$i) \lim_{j \rightarrow -\infty} P_j f = 0.$$

$$ii) \lim_{j \rightarrow \infty} P_j f = f.$$

Demostración. *i)* Puesto que $\|P_j\| = 1$, basta probar el resultado para funciones en $L_2(\mathbb{R})$ con soporte compacto. Si f tiene soporte en $[-a, a]$, entonces al aplicar las desigualdades de Cauchy-Schwarz y de Minkowski

se tiene

$$\begin{aligned}
 \|P_j f\|_2^2 &= \left\| \sum_{k \in \mathbb{Z}} \langle f, \varphi_{jk} \rangle \varphi_{jk} \right\|_2^2 \\
 &= \sum_{k \in \mathbb{Z}} |\langle f, \varphi_{jk} \rangle|^2 \\
 &= \sum_{k \in \mathbb{Z}} \left| \int_{-a}^a f(t) 2^{j/2} \varphi(2^j t - k) dx \right|^2 \\
 &\leq \sum_{k \in \mathbb{Z}} \left(\int_{-a}^a |f(t)|^2 dt \right) 2^{j/2} \left(\int_{-a}^a |\varphi(2^j t - k)| dt \right)^2 \\
 &= \|f\|_2^2 \sum_{k \in \mathbb{Z}} \left(\int_{-2^j a - k}^{2^j a - k} |\varphi(z)| dz \right)^2.
 \end{aligned}$$

Si $2^j a < 1/2$, entonces estas integrales están definidas sobre intervalos ajenos cuya unión se escribe $\Omega_j = \cup_{k \in \mathbb{Z}} (-2^j a - k, 2^j a - k)$, con $\cap_j \Omega_j = \mathbb{Z}$, el cual tiene medida cero. Por tanto,

$$\|P_j f\|_2^2 \leq \|f\|_2^2 \int_{\Omega_j} |\varphi(z)|^2 dz \rightarrow 0, \quad j \rightarrow -\infty$$

por el teorema de la convergencia dominada de Lebesgue. $\square$

2.4. Ecuación de escala

Puesto que el conjunto $\{\varphi(\cdot - k)\}_{k \in \mathbb{Z}}$ constituye una base ortonormal de V_0 entonces cada $f \in V_0$ se puede expresar como

$$f = \sum_{n \in \mathbb{Z}} a_n \varphi_{0n}, \quad \varphi_{0n}(x) = \varphi(x - n).$$

Ahora, como $\varphi \in V_0$, y $V_0 \subset V_1$, se tiene entonces $\varphi \in V_1$. Pero la propiedad de dilatación implica que $\varphi(2^{-1}\cdot) \in V_0$. En consecuencia, se puede expandir

$$\varphi(2^{-1}t) = \sum_{n \in \mathbb{Z}} g_n \varphi(t - n), \quad t \in \mathbb{R},$$

para algunos coeficientes $(g_n)_{n \in \mathbb{Z}}$. O de manera equivalente,

$$\varphi(t) = \sum_{n \in \mathbb{Z}} g_n \varphi(2t - n), \quad t \in \mathbb{R}, \quad (2.4.1)$$

en donde las constantes de estructura (los (g_n)) satisfacen $\sum_{n \in \mathbb{Z}} |g_n|^2 < \infty$. La relación (2.4.1) se llama ecuación de escala. Los coeficientes g_n constituyen un filtro $\mathbf{g} = (g_n)_{n \in \mathbb{Z}}$ asociado a la función de escala.

Ejemplo 2.4.1. Si $\varphi(t) = \chi_{[0,1)}(t)$, entonces claramente

$$\varphi(t) = \varphi(2t) + \varphi(2t - 1)$$

es la ecuación de escala, con las constantes de estructura $g_0 = 1$, $g_1 = 1$ y $g_n = 0$, en otro caso.

A continuación se dan algunas propiedades de las constantes de estructura.

Proposición 2.4.1. *Los coeficientes de la ecuación de escala satisfacen las siguientes propiedades:*

$$g_n = 2 \int_{\mathbb{R}} \varphi(t) \overline{\varphi(2t - n)} dt, \quad n \in \mathbb{Z} \quad (2.4.2)$$

$$\sum_{k \in \mathbb{Z}} g_k \bar{g}_{2n+k} = 2\delta_{0n} \quad (\text{delta de Kronecker}). \quad (2.4.3)$$

$$\sum_{n \in \mathbb{Z}} |g_n|^2 = 2 \quad (2.4.4)$$

Si también $\varphi \in L_1(\mathbb{R})$, $\int_{\mathbb{R}} \varphi \neq 0$ y la ecuación (2.4.1) converge en $L_1(\mathbb{R})$, entonces

$$\sum_{n \in \mathbb{Z}} g_n = 2. \quad (2.4.5)$$

Demostración. Puesto que $g_n/\sqrt{2}$ son los coeficientes de Fourier de $\varphi \in V_1$ con respecto a la base ortonormal $\sqrt{2}\varphi(2t - n)$, se tiene

$$\frac{g_n}{\sqrt{2}} = \int_{\mathbb{R}} \varphi(t) \sqrt{2} \overline{\varphi(2t - n)} dt,$$

o lo que es lo mismo,

$$g_n = 2 \int_{\mathbb{R}} \varphi(t) \overline{\varphi(2t - n)} dt.$$

De la propiedad 5 de la definición de AMR se tiene

$$\int_{\mathbb{R}} \varphi(t-n) \overline{\varphi(t)} dt = \delta_{0n}.$$

Al sustituir (2.4.1) y aplicar la identidad de Parseval y la ortogonalidad se tiene

$$\begin{aligned} \delta_{0n} &= \sum_{k,m \in \mathbb{Z}} g_k \bar{g}_m \int_{\mathbb{R}} \varphi(2t-2n-k) \overline{\varphi(2t-m)} dt \\ &= \frac{1}{2} \sum_{2n+k=m} g_k \bar{g}_m, \end{aligned}$$

lo cual es lo mismo que (2.4.3). En particular, si se toma $n = 0$ en la última expresión, se obtiene

$$\sum_{k=m} g_k \bar{g}_k = \sum_{k \in \mathbb{Z}} |g_k|^2 = 2.$$

Si, además, se tiene $\varphi \in L_1(\mathbb{R})$ con $\int_{\mathbb{R}} \varphi(t) dt \neq 0$, entonces al integrar (2.4.1) término a término se llega

$$\begin{aligned} \int_{\mathbb{R}} \varphi(t) dt &= \sum_{n \in \mathbb{Z}} g_n \int_{\mathbb{R}} \varphi(2t-n) dt \\ &= \frac{1}{2} \sum_{n \in \mathbb{Z}} g_n \int_{\mathbb{R}} \varphi(t) dt, \end{aligned}$$

al dividir por $\int_{\mathbb{R}} \varphi$ se obtiene $\sum_{n \in \mathbb{Z}} g_n = 2$. □

2.5. Construcción de la función de escala

Para construir la función de escala, es necesario encontrar los coeficientes g_n . Una forma de hacerlo es vía la transformada de Fourier, puesto que de manera directa es difícil (ver p. e., [25], [36]). En consecuencia, al aplicar la transformada de Fourier a la ecuación (2.4.1) se obtiene

$$\begin{aligned} \hat{\varphi}(\omega) &= \frac{1}{2} \sum_{n \in \mathbb{Z}} g_n e^{-in\omega/2} \hat{\varphi}\left(\frac{\omega}{2}\right) \\ &= \hat{\varphi}\left(\frac{\omega}{2}\right) P(e^{-i\omega/2}) \end{aligned} \tag{2.5.1}$$

donde el polinomio P es dado por

$$P(z) = \frac{1}{2} \sum_{n \in \mathbb{Z}} g_n z^n.$$

Al iterar (2.5.1) se tiene

$$\begin{aligned} \hat{\varphi}(\omega) &= P(e^{-i\omega/2}) \hat{\varphi}\left(\frac{\omega}{2}\right) \\ \hat{\varphi}\left(\frac{\omega}{2}\right) &= P(e^{-i\omega/2^2}) \hat{\varphi}\left(\frac{\omega}{2^2}\right) \end{aligned}$$

$$\hat{\varphi}(\omega) = P(e^{-i\omega/2}) P(e^{-i\omega/2^2}) \hat{\varphi}\left(\frac{\omega}{2^2}\right).$$

Continuando de esta manera se obtiene

$$\begin{aligned} \hat{\varphi}(\omega) &= P(e^{-i\omega/2}) \dots P(e^{-i\omega/2^n}) \hat{\varphi}\left(\frac{\omega}{2^n}\right) \\ &= \left(\prod_{j=1}^n P(e^{-i\omega/2^j}) \right) \hat{\varphi}\left(\frac{\omega}{2^n}\right). \end{aligned}$$

Para una función de escala dada φ , la ecuación precedente se tiene para cada n . En el límite, cuando $n \rightarrow \infty$, la última ecuación se transforma

$$\hat{\varphi}(\omega) = \left(\prod_{j=1}^{\infty} P(e^{-i\omega/2^j}) \right) \hat{\varphi}(0).$$

Si φ satisface la condición de normalización $\int_{\mathbb{R}} \varphi = 1$, entonces $\hat{\varphi}(0) = 1$ y así,

$$\hat{\varphi}(\omega) = \prod_{j=1}^{\infty} P(e^{-i\omega/2^j}). \quad (2.5.2)$$

Por tanto, si el producto $\prod_{j=1}^{\infty} P(e^{-i\omega/2^j})$ converge, entonces la función de escala queda determinada salvo un factor no nulo $\hat{\varphi}(0)$, que es su media. En consecuencia, la única función de escala asociada al filtro ***g*** está dada por (2.5.2). Es decir, si la función P asociada al filtro ***g*** cumple cierta propiedad de convergencia, entonces se tiene $\hat{\varphi}$ y, antitransformando, se obtiene φ . En resumen, se tiene

Proposición 2.5.1. Sea $\mathbf{g}$ un filtro y P el polinomio dado por

$$P(z) = \frac{1}{2} \sum_{n \in \mathbb{Z}} g_n z^n.$$

Si la función Φ definida por

$$\Phi(\omega) = \lim_{n \rightarrow \infty} \prod_{j=1}^n P(e^{-i\omega/2^j})$$

está en $L_2(\mathbb{R})$ y verifica $\lim_{|\omega| \rightarrow \infty} \Phi(\omega) = 0$. Entonces existe una función de escala φ asociada al filtro $\mathbf{g}$ y determinada por $\hat{\varphi} = \Phi$ con $\int \varphi = 1$.

La siguiente proposición permite dar condiciones sobre la ortonormalidad de la base $\{\varphi_{0k}\}_{k \in \mathbb{Z}}$ en términos de los coeficientes g_k .

Proposición 2.5.2. El sistema $\{\varphi_{0k}\}_{k \in \mathbb{Z}}$ es ortonormal si y sólo si la transformada de Fourier de φ satisface

$$2\pi \sum_{k \in \mathbb{Z}} |\hat{\varphi}(\omega + 2k\pi)|^2 = 1.$$

Demostración. Como $\varphi(t - k)$ forma una base ortonormal en V_0 , entonces al aplicar el teorema de Plancherel (ver p.e., [54]) se tiene

$$\begin{aligned} \delta_{0m} &= \int_{\mathbb{R}} \varphi(t) \overline{\varphi(t - m)} dt \\ &= \int_{\mathbb{R}} \hat{\varphi}(\omega) \overline{\hat{\varphi}(\omega)} e^{-im\omega} d\omega \\ &= \int_{\mathbb{R}} |\hat{\varphi}(\omega)|^2 e^{im\omega} d\omega \\ &= \sum_{k \in \mathbb{Z}} \int_{2\pi k}^{2\pi(k+1)} |\hat{\varphi}(\omega)|^2 e^{im\omega} d\omega \\ &= \int_0^{2\pi} e^{im\omega} \left(\sum_{k \in \mathbb{Z}} |\hat{\varphi}(\omega + 2k\pi)|^2 \right) d\omega. \end{aligned}$$

Sea $F(\omega) = 2\pi \sum_{k \in \mathbb{Z}} |\hat{\varphi}(\omega + 2k\pi)|^2$, entonces

$$\frac{1}{2\pi} \int_0^{2\pi} e^{im\omega} F(\omega) d\omega = \delta_{0m}. \quad (*)$$

La función F es 2π -periódica ya que

$$\begin{aligned} F(\omega + 2\pi) &= 2\pi \sum_{k \in \mathbb{Z}} |\hat{\varphi}(\omega + 2\pi(k+1))|^2 \\ &= 2\pi \sum_{j \in \mathbb{Z}} |\hat{\varphi}(\omega + 2\pi j)|^2 = F(\omega). \end{aligned}$$

Como F es periódica, su serie de Fourier, $\sum_{m \in \mathbb{Z}} c_m e^{imt}$, donde los coeficientes de Fourier son dados por $c_m = \frac{1}{2\pi} \int_0^{2\pi} F(\omega) e^{-im\omega} d\omega$. Por tanto, la condición de ortonormalidad (*), es equivalente a $c_{-m} = \delta_{m0}$, lo cual a su vez es equivalente a $F(\omega) = 1$. $\square$

Como consecuencia de este resultado se tiene la siguiente condición necesaria sobre el polinomio $P(z)$ para la existencia de un AMR.

Corolario 2.5.1. *El polinomio $P(z) = \sum_{n \in \mathbb{Z}} g_n z^n$ satisface*

$$|P(e^{-it})|^2 + |P(e^{-i(t+\pi)})|^2 = 1, \quad 0 \leq t \leq 2\pi. \quad (2.5.3)$$

Demostración. De los resultados anteriores se tiene

$$\sum_{n \in \mathbb{Z}} |\hat{\varphi}(\omega + 2n\pi)|^2 = \frac{1}{2\pi} \quad \text{y} \quad \hat{\varphi}(\omega) = P(e^{-i\omega/2}) \hat{\varphi}\left(\frac{\omega}{2}\right).$$

$$\begin{aligned} \frac{1}{2\pi} &= \sum_{n \in \mathbb{Z}} |\hat{\varphi}(\omega + 2n\pi)|^2 = \sum_{n \text{ par}} + \sum_{n \text{ impar}} \\ &= \sum_{k \in \mathbb{Z}} |\hat{\varphi}(\omega + (2k)2\pi)|^2 + \sum_{k \in \mathbb{Z}} |\hat{\varphi}(\omega + (2k+1)2\pi)|^2 \\ &= \sum_{k \in \mathbb{Z}} \left(|P(e^{-i(\frac{\omega}{2} + 2k\pi)})|^2 \left| \hat{\varphi}\left(\frac{\omega}{2} + 2k\pi\right) \right|^2 \right. \\ &\quad \left. + |P(e^{-i(\frac{\omega}{2} + (2k+1)\pi)})|^2 \left| \hat{\varphi}\left(\frac{\omega}{2} + (2k+1)\pi\right) \right|^2 \right) \\ &= |P(e^{-i\frac{\omega}{2}})|^2 \sum_{k \in \mathbb{Z}} \left| \hat{\varphi}\left(\frac{\omega}{2} + 2k\pi\right) \right|^2 \\ &\quad + |P(-e^{-i\frac{\omega}{2}})|^2 \sum_{k \in \mathbb{Z}} \left| \hat{\varphi}\left(\left(\frac{\omega}{2} + \pi\right) + 2k\pi\right) \right|^2 \\ &= |P(e^{-i\frac{\omega}{2}})|^2 \frac{1}{2\pi} + |P(-e^{-i\frac{\omega}{2}})|^2 \frac{1}{2\pi}. \end{aligned}$$

En consecuencia,

$$1 = \left| P(e^{-i\frac{\omega}{2}}) \right|^2 + \left| P(-e^{-i\frac{\omega}{2}}) \right|^2.$$

□

Corolario 2.5.2. *En términos de g_k , (2.5.3) se transforma en*

$$\sum_{n \in \mathbb{Z}} g_{n-2k} \bar{g}_{n-2m} = 2\delta_{k-m,0}.$$

Demostración. Observe que (2.5.3) en términos de los coeficientes g_n está dado por

$$\frac{1}{4} \sum_{n,k \in \mathbb{Z}} g_n \bar{g}_k e^{-i(n-k)t} + \frac{1}{4} \sum_{n,k \in \mathbb{Z}} g_n \bar{g}_k (-1)^{n-k} e^{-i(n-k)t} = 1.$$

Los términos impares se cancelan y por lo tanto se tiene

$$\sum_{n-k \text{ par}} g_n \bar{g}_k e^{-i(n-k)t} = \sum_{j \in \mathbb{Z}} c_j e^{-2ijt} = 2$$

donde

$$c_j = \sum_{n-k=2j} g_n \bar{g}_k.$$

Esto es válido para todo t , luego $c_j = 2\delta_j$. Por tanto

$$\sum_{n \in \mathbb{Z}} g_n \bar{g}_{n-2j} = 2\delta_j.$$

De hecho, esto es equivalente con la afirmación a probar. □

2.6. Descomposición y reconstrucción

En esta sección se describirán algoritmos de descomposición y reconstrucción asociados a un AMR. Estos algoritmos se utilizarán junto con el análisis multirresolución en la descomposición y reconstrucción de señales en donde tanto la función de escala como la wavelet son funciones continuas.

2.6.1. Algoritmo de descomposición

Sean $c_{j,k}$ y $d_{j,k}$ los coeficientes de la función de escala φ y de la wavelet ψ , respectivamente, para $j, k \in \mathbb{Z}$, definidos por

$$c_{j,k} := \int_{\mathbb{R}} f(t) \varphi_{jk}(t) dt \quad (2.6.1)$$

$$d_{j,k} := \int_{\mathbb{R}} f(t) \psi_{jk}(t) dt, \quad (2.6.2)$$

donde

$$\varphi_{jk}(t) = 2^{j/2} \varphi(2^j t - k), \quad \psi_{jk}(t) = 2^{j/2} \psi(2^j t - k).$$

φ y ψ son, respectivamente, la función de escala y la wavelet madre.

Ahora bien, como $\varphi_{jk}(t) = 2^{j/2} \varphi(2^j t - k)$, entonces existe h_m tal que

$$\begin{aligned} \varphi_{jk}(t) &= \sum_{m \in \mathbb{Z}} h_m 2^{j/2} \varphi_{1m}(2^j t - k) \\ &= \sum_{m \in \mathbb{Z}} h_m 2^{(j+1)/2} \varphi(2^{j+1} t - 2k - m) \\ &= \sum_{m \in \mathbb{Z}} h_m \varphi_{j+1, m+2k}(t) \\ &= \sum_{m \in \mathbb{Z}} h_{m-2k} \varphi_{j+1, m}(t). \end{aligned} \quad (2.6.3)$$

Reemplazando este valor en (2.6.1), se obtiene

$$\begin{aligned} c_{j,k} &= \int_{\mathbb{R}} f(t) \sum_{m \in \mathbb{Z}} h_{m-2k} \varphi_{j+1, m}(t) dt \\ &= \sum_{m \in \mathbb{Z}} h_{m-2k} \int_{\mathbb{R}} f(t) \varphi_{j+1, m}(t) dt \\ &= \sum_{m \in \mathbb{Z}} h_{m-2k} c_{j+1, m}, \end{aligned}$$

por tanto,

$$c_{j,k} = \sum_{m \in \mathbb{Z}} h_{m-2k} c_{j+1, m}. \quad (2.6.4)$$

Como $V_0 \subset V_1$, para cada $\varphi \in V_0$ también se satisface $\varphi \in V_1$. Además, $\{\varphi_{1k}, k \in \mathbb{Z}\}$ es una base ortonormal para V_1 , entonces existe una sucesión $(h_k) \in \ell_2(\mathbb{Z})$ tal que

$$\varphi(t) = \sum_{k \in \mathbb{Z}} h_k \varphi_{1k}(t), \quad (2.6.5)$$

por tanto, los elementos de la sucesión se puede escribir como

$$h_k = \langle \varphi, \varphi_{1k} \rangle \quad \text{y} \quad (h_k) \in \ell_2.$$

La ecuación (2.6.5) relaciona funciones con diferentes factores de escala, también se conoce como ecuación de dilatación. Para la base de Haar se tiene

$$h_k = \begin{cases} 1/\sqrt{2}, & k = 0, 1 \\ 0, & \text{en otro caso.} \end{cases}$$

Si φ es la función de escala de un AMR, entonces la wavelet madre ψ se relaciona con φ por medio de la ecuación

$$\psi(t) = \sum_{k \in \mathbb{Z}} (-1)^k h_{1-k} \varphi_{1k}(t). \quad (2.6.6)$$

Al sustituir (2.6.6) en (2.6.2) se obtiene

$$d_{j,k} = \sum_{p \in \mathbb{Z}} (-1)^p h_{1-p+2k} c_{j+1,p}. \quad (2.6.7)$$

Si los coeficientes de escala en cualquier nivel j son dados, entonces todos los coeficientes de la función escala de nivel inferior para $J < j$, se pueden calcular recursivamente usando la ecuación (2.6.4), mientras que todos los coeficientes wavelet de nivel inferior ($J < j$) se calculan aplicando (2.6.7).

Si $c_{j,\cdot}$ y $d_{j,\cdot}$ representan los coeficientes de la función de escala y wavelet en el nivel j , respectivamente, la Figura 2.6.1 representa el algoritmo de descomposición en forma esquemática. Por ejemplo, la flecha que relaciona los coeficientes c_{j-1} y c_{j-2} , indica que c_{j-2} se calcula sólo usando c_{j-1} .

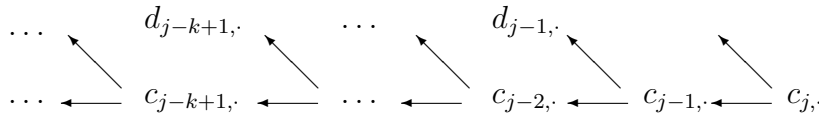


Figura 2.6.1

Observe que las fórmulas (2.6.4) y (2.6.7) comparten un hecho interesante, esto es, en cada ecuación, si el índice de dilatación k se incrementa en uno, todos los índices de (h_m) se desplazan en dos unidades; lo cual significa que si existe solamente un número finito de términos no nulos en la sucesión (h_m) , entonces aplicando el algoritmo de descomposición a un conjunto de coeficientes de escala no nulos en el nivel $j + 1$, se obtendrá sólo la mitad de coeficientes no nulos en el nivel j . Este proceso en teoría de señales se conoce como *downsampling*. Un resultado análogo se tiene para los coeficientes wavelet.

Para expresar lo anterior en la terminología de filtros, recuerde que la convolución de dos sucesiones en $\ell_2(\mathbb{Z})$

$$x = (\dots, x_{-1}, x_0, x_1, \dots) \quad \text{y} \quad y = (\dots, y_{-1}, y_0, y_1, \dots)$$

se define por

$$(x * y)_m := \sum_{k \in \mathbb{Z}} x_k y_{m-k}.$$

En consecuencia, (2.6.4) se puede expresar como

$$c_{j-1,k} = \sum_{m \in \mathbb{Z}} \tilde{h}_{2k-m} c_{j,m} = (\tilde{h} * c_j)_{2k}, \quad (2.6.8)$$

note que se reemplazó j por $j - 1$ y para simplificar se utilizó la notación $\tilde{y}_m = y_{-m}$. Si se define el operador downsampling para la sucesión x como

$$((\downarrow 2)x)_k := x_{2k}, \quad k \in \mathbb{Z},$$

entonces (2.6.8) se puede escribir

$$c_{j-1,\cdot} = (\downarrow 2)(\tilde{h} * c_j). \quad (2.6.9)$$

De un procedimiento similar se obtiene, con $g_m = (-1)^m h_{1-m}$,

$$d_{j-1,\cdot} = (\downarrow 2)(\tilde{g} * c_j). \quad (2.6.10)$$

Al algoritmo de descomposición, Mallat lo llamó *algoritmo piramidal* [47], mientras Daubechies lo llamó *algoritmo de cascada* [25].

2.6.2. Algoritmo de reconstrucción

Recuerde que dado un AMR, el conjunto de funciones linealmente independientes φ_{jk} que generan a V_j son las funciones de escala, mientras que el conjunto de funciones linealmente independientes ψ_{jk} que generan a W_j son las wavelets. En otras palabras, $\{\varphi_{jk}\}_{k \in \mathbb{Z}}$ y $\{\psi_{jk}\}_{k \in \mathbb{Z}}$ son generadas, respectivamente, por φ y ψ , esto es,

$$\varphi_{jk}(t) = 2^{j/2} \varphi(2^j t - k) \quad \text{y} \quad \psi_{jk}(t) = 2^{j/2} \psi(2^j t - k), \quad \forall k \in \mathbb{Z}$$

forman las bases ortonormales para V_j y W_j , respectivamente. Definiendo

$$a_{2k} = \langle \varphi_{10}, \varphi_{0k} \rangle, \quad a_{2k-1} = \langle \varphi_{11}, \varphi_{0k} \rangle$$

$$b_{2k} = \langle \varphi_{10}, \psi_{0k} \rangle, \quad b_{2k-1} = \langle \varphi_{11}, \psi_{0k} \rangle,$$

donde $a_k = h_{-k}$ y $b_k = (-1)^k h_{k+1}$. Entonces

$$\begin{aligned} c_{j,k} &= \sum_{m \in \mathbb{Z}} a_{2m-k} c_{j-1,m} + b_{2m-k} d_{j-1,m} \\ &= \sum_{m \in \mathbb{Z}} h_{k-2m} c_{j-1,m} + (-1)^k h_{2m-k+1} d_{j-1,m}. \end{aligned} \quad (2.6.11)$$

Observe que esta última expresión es casi la suma de dos convoluciones. La diferencia está en que el índice para la convolución es $k - m$ mientras acá aparece $k - 2m$. En otras palabras, (2.6.11) es una convolución pero sin los términos impares (falta $h_{k-(2m-1)}$). Para que (2.6.11) sea una convolución, se altera la sucesión original intercalando ceros entre sus componentes y obteniendo una nueva sucesión que contiene ceros en todas sus entradas impares. Este procedimiento se llama *upsampling*, denotado por $(\uparrow 2)$. Más explícitamente, si $x = (\dots, x_{-2}, x_{-1}, x_0, x_1, x_2, \dots)$, entonces

$$((\uparrow 2)x)_k = (\dots, x_{-2}, 0, x_{-1}, 0, x_0, 0, x_1, 0, x_2, \dots)$$

o de manera equivalente,

$$((\uparrow 2)x)_k = \begin{cases} x_{k/2}, & \text{si } k \text{ es par,} \\ 0, & \text{si } k \text{ es impar.} \end{cases}$$

En consecuencia,

$$c_{j,k} = (((\uparrow 2)c_{j-1}) * h)_k + (((\uparrow 2)d_{j-1}) * g)_k. \quad (2.6.12)$$

La Figura 2.6.2 representa el algoritmo de reconstrucción en forma esquemática.

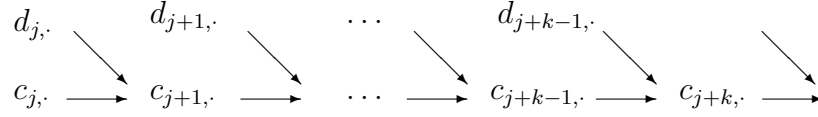


Figura 2.6.2

2.7. Wavelet biortogonales

A partir de las condiciones de ortogonalización de la función de escala $\varphi(t)$ y el hecho de que la función wavelet $\psi(t)$ surge de los complementos ortogonales de los espacios V_j se pudo establecer la descomposición y reconstrucción de una señal al construirse cuatro filtros estrechamente relacionados $h, l, \tilde{h}, \tilde{l}$ estableciéndose una reconstrucción perfecta de la señal. Se pueden establecer condiciones para los filtros finitos $h, l, \tilde{h}, \tilde{l}$ que permiten también una reconstrucción perfecta sin la condición de tomar los complementos ortogonales dando lugar a las wavelet biortogonales. El análisis multirresolución con el que se determinan estas wavelet se define a partir de dos espacios de aproximación

$$\dots V_2 \subset V_1 \subset V_0 \subset V_{-1} \subset V_{-2} \dots$$

$$\dots \tilde{V}_2 \subset \tilde{V}_1 \subset \tilde{V}_0 \subset \tilde{V}_{-1} \subset \tilde{V}_{-2} \dots,$$

donde los W_j y $\tilde{W}_j$ son los complementos no ortogonales de V_j en V_{j-1} y de $\tilde{V}_j$ en $\tilde{V}_{j-1}$ respectivamente, con $\tilde{W}_j \perp V_j$ y $W_j \perp \tilde{V}_j$. Se establecen los coeficientes para la descomposición y reconstrucción de una señal $c^0 = (c_n^0)_{n \in \mathbb{N}}$. Descomposición con filtros h y g

$$c_n^1 = \sum_k h_{2n-k} c_k^0 \quad \text{y} \quad d_n^1 = \sum_k g_{2n-k+1} c_k^0. \quad (2.7.1)$$

Reconstrucción con filtros $\tilde{h}$ y $\tilde{g}$

$$\tilde{c}_l^0 = \sum_n \left[\tilde{h}_{2n-l} c_n^1 + \tilde{g}_{2n-l+1} d_n^1 \right]. \quad (2.7.2)$$

Y se determinan las condiciones para la reconstrucción perfecta $\tilde{c}_0 = c_0$

$$\sum_n \left[\tilde{h}_{2n-l} h_{2n-k} + \tilde{g}_{2n-l+1} g_{2n-k+1} \right] = \delta_{lk}$$

la cual es equivalente a

$$\frac{1}{2} \left[h(z) \bar{\tilde{h}}(z) - g(z) \bar{\tilde{g}}(z) \right] = 1 \quad (2.7.3)$$

$$\frac{1}{2} \left[h(-z) \bar{\tilde{h}}(z) - g(-z) \bar{\tilde{g}}(z) \right] = 0,$$

donde $h(z) = \sum_n h_n z^n$ y $\bar{a} = \sum_n a_{-n} z^n = \sum_n a_n z^{-n}$ para $|z| = 1$. De donde $\bar{\tilde{h}}(z)$ y $\bar{\tilde{g}}(z)$ se pueden expresar como producto de polinomios en z

$$\bar{\tilde{h}}(z) = g(-z) p(z) \quad \text{y} \quad \bar{\tilde{g}}(z) = h(-z) p(z).$$

($h(-z)$ y $g(-z)$ no tiene ceros en común). Luego de (2.7.3) se obtiene

$$p(z) [h(z) g(-z) + h(-z) g(z)] = 2.$$

Esto sólo es posible si

$$p(z) = \alpha z^k \quad (2.7.4)$$

$$h(z) g(-z) + h(-z) g(z) = 2\alpha^{-1} z^{-k}$$

para $\alpha \in \mathbb{C} - \{0\}$, $k \in \mathbb{Z}$. Luego

$$\bar{\tilde{h}}(z) = \alpha z^k g(-z), \quad \bar{\tilde{g}}(z) = \alpha z^k h(-z) \quad (2.7.5)$$

las condiciones (2.7.4), (2.7.5) son suficientes y necesarias para establecer la descomposición y reconstrucción de una señal a partir de (2.7.1), (2.7.2). Tomando $k = 0$ y $\alpha = -1$ se tiene

$$g_n = (-1)^{n+1} \tilde{h}_{-n}, \quad \tilde{g} = (-1)^{n+1} h_{-n}.$$

Además las condiciones (2.4.3) y (2.4.5) se transforma en

$$h(z) \bar{\tilde{h}}(z) + h(-z) \bar{\tilde{h}}(-z) = 2$$

$$\sum_k h_n \tilde{h}_{n+2k} = \delta_{k0}.$$

Si $h = \tilde{h}$, se obtiene las condiciones establecidas en el caso de las wavelet ortogonales. Para determinar los filtros h y $\tilde{h}$ nos trasladaremos al dominio de la frecuencia considerando las ecuaciones de dilatación en este dominio

$$\varphi(x) = \sqrt{2} \sum_n h_n \varphi(2x - n) \longrightarrow \widehat{\varphi}(\xi) = (2\pi)^{-1/2} \prod_{j=1}^{\infty} m_0(2^j \xi)$$

$$\tilde{\varphi}(x) = \sqrt{2} \sum_n \tilde{h}_n \tilde{\varphi}(2x - n) \longrightarrow \widehat{\tilde{\varphi}}(\xi) = (2\pi)^{-1/2} \prod_{j=1}^{\infty} \tilde{m}_0(2^j \xi)$$

con

$$m_0(2^j \xi) = 2^{-1/2} \sum_n h_n e^{-in\xi} \quad y \quad \tilde{m}_0(2^j \xi) = 2^{-1/2} \sum_n \tilde{h}_n e^{-in\xi}$$

donde $m_0(0) = 1 = \tilde{m}_0(0)$ y $\sum_n h_n = h(1) = \sqrt{2} = \tilde{h}(1) = \sum_n \tilde{h}_n$.

Además,

$$\psi(x) = \sqrt{2} \sum_n g_{n+1} \varphi(2x - n) = \sqrt{2} \sum_n (-1)^n \tilde{h}_{-n-1} \varphi(2x - n)$$

$\psi(x) = \sqrt{2} \sum_n \tilde{g}_{n+1} \tilde{\varphi}(2x - n) = \sqrt{2} \sum_n (-1)^n h_{-n-1} \tilde{\varphi}(2x - n)$. En el dominio de la frecuencia

$$\widehat{\psi}(x) = e^{i\xi/2} \overline{\tilde{m}_0\left(\frac{\xi}{2} + \pi\right)} \widehat{\varphi}(\xi/2)$$

$$\widehat{\tilde{\psi}}(x) = e^{i\xi/2} \overline{m_0\left(\frac{\xi}{2} + \pi\right)} \widehat{\tilde{\varphi}}(\xi/2),$$

esta función wavelet generaran una base Riesz si $m_0(\pi) = 0 = \tilde{m}_0(\pi)$

$$\sum_n (-1)^n h_n = h(-1) = 0 = \tilde{h}(-1) = \sum_n (-1)^n \tilde{h}_n,$$

es decir, si $\psi_{jk}(x) = 2^{-j/2} \psi(2^{-j}x - k)$

$$\tilde{\psi}_{jk}(x) = 2^{-j/2} \tilde{\psi}(2^{-j}x - k)$$

entonces

$$f = \sum_{j,k \in \mathbb{Z}} \langle f, \tilde{\psi}_{jk} \rangle \psi_{jk} = \sum_{j,k \in \mathbb{Z}} \langle f, \psi_{jk} \rangle \tilde{\psi}_{jk}$$

para $f \in L^2(\mathbb{R})$. El filtro h_n está determinado por los coeficientes del polinomio $m_0(\xi)$ que a su vez establece la función escala φ la cual es simétrica si $m_0(-\xi) = m_0(\xi)$ ó $m_0(\xi) = \text{polinomio en } \cos \xi$.

La función Haar, φ_{Haar} , es simétrica con respecto $x = 1/2$, $\varphi_{Haar}(1-x) = \varphi_{Haar}(x)$ le corresponde un polinomio trigonométrico $m_0(\xi)$ que satisface $m_0(-\xi) = e^{i\xi} m_0(\xi)$ ó $m_0(\xi) = e^{-i\xi/2} \cos(\xi/2)$ polinomio en $\cos \xi$. La condición para una reconstrucción perfecta está dada por $h(z) \tilde{h}(z) + h(-z) \tilde{h}(-z) = 2$ o equivalentemente

$$m_0(\xi) \overline{\tilde{m}_0(\xi)} + m_0(\xi + \pi) \overline{\tilde{m}_0(\xi + \pi)} = 1. \quad (2.7.6)$$

Así, si $\tilde{m}_0(\xi)$ es solución de (2.7.6) para un m_0 fijo tal que $m_0(-\xi) = m_0(\xi)$ entonces $\tilde{m}'_0 = \frac{1}{2}[\tilde{m}_0(\xi) + \tilde{m}_0(-\xi)]$ es también solución de (2.7.6) que satisface $\tilde{m}'_0(\xi) = \tilde{m}'_0(-\xi)$. La siguiente proposición resume algunas propiedades útiles en sistemas biortogonales

Proposición 2.7.1. *Sea $m_0(\xi)$ un polinomio trigonométrico con coeficientes reales*

- a) (i) Si $m_0(-\xi) = m_0(\xi)$ entonces $m_0(\xi) = (\cos \xi/2)^{2l} P_0(\cos \xi)$.
 (ii) Si $m_0(-\xi) = e^{-i\xi} m_0(\xi)$ entonces $m_0(\xi) = e^{-i\xi/2} (\cos \xi/2)^{2l+1} P_0(\cos \xi)$, donde P_0 es un polinomio tal que $P_0(-1) \neq 0$ y $l \in \mathbb{N}$.
- b) Si $\tilde{m}_0$ es solución de (2.7.6) entonces $\tilde{m}_0(\xi) = (\cos \xi/2)^{2\tilde{l}} \tilde{P}_0(\cos \xi)$ (en el caso a. i) ó $\tilde{m}_0(\xi) = e^{-i\xi/2} (\cos \xi/2)^{2\tilde{l}+1} \tilde{P}_0(\cos \xi)$ (en el caso a. ii) con $\tilde{P}_0(-1) \neq 0$ y $\tilde{l} \in \mathbb{N}$, donde P_0 y $\tilde{P}_0$ son construidos por

$$\tilde{P}_0(\cos \xi) \tilde{P}_0(\cos \xi) = \sum_{n=0}^{k-1} \binom{k-1+n}{n} \left(\frac{\sin^2 \xi}{2} \right)^n + \left(\frac{\sin^2 \xi}{2} \right)^k R(\cos \xi)$$

con $k = l + \tilde{l}$ y R un polinomio impar.

Las funciones ${}_N\varphi$ $B - spline$ de orden N proporcionan polinomios simétricos m_0

$$B - spline \text{ constante } {}_1\varphi(x) = \begin{cases} 1, & 0 \leq x \leq 1 \\ 0, & \text{en otra parte} \end{cases}$$

$$B - spline \text{ lineal } {}_2\varphi(x) = \begin{cases} 1+x, & -1 \leq x \leq 0 \\ 1-x, & 0 \leq x \leq 1 \\ 0 & \text{en otra parte} \end{cases}$$

$$B - spline \text{ cuadrática } {}_3\varphi(x) = \begin{cases} (1+x)^2/2, & -1 \leq x \leq 0 \\ -(x-1/2)^2+3/4, & 0 \leq x \leq 1 \\ (x-2)^2/2 & 1 \leq x \leq 2 \\ 0 & \text{en otra parte.} \end{cases}$$

Se verifica que ${}_{2L}\phi(-x) = {}_{2L}\phi(x)$ y ${}_{2L+1}\phi(1-x) = {}_{2L+1}\phi(x)$. El polinomio ${}_Nm_0$ correspondiente a ${}_N\phi(x)$ está dado por

$$\begin{aligned} {}_Nm_0(\xi) &= \left(\frac{1+e^{-i\xi}}{2} \right) N e^{i\xi \lfloor N/2 \rfloor} = e^{-ik\xi/2} \left(\frac{\cos \xi}{2} \right)^N \\ &= \sum_{n=-\lfloor N/2 \rfloor}^{n=\lfloor N/2 \rfloor} 2^{-N} \binom{N}{n + \lfloor N/2 \rfloor} e^{-in\xi}, \end{aligned}$$

donde ${}_{2L}m_0(-\xi) = {}_{2L}m_0(\xi)$ y ${}_{2L+1}m_0(-\xi) = e^{i\xi} {}_{2L+1}m_0(\xi)$ (tomando $l = L$ y $P_0 = 1$). El polinomio simétrico ${}_{N,\tilde{N}}m_0(\xi)$ asociado a ${}_Nm_0(\xi)$ que satisface la condición (2.7.6) está dado por

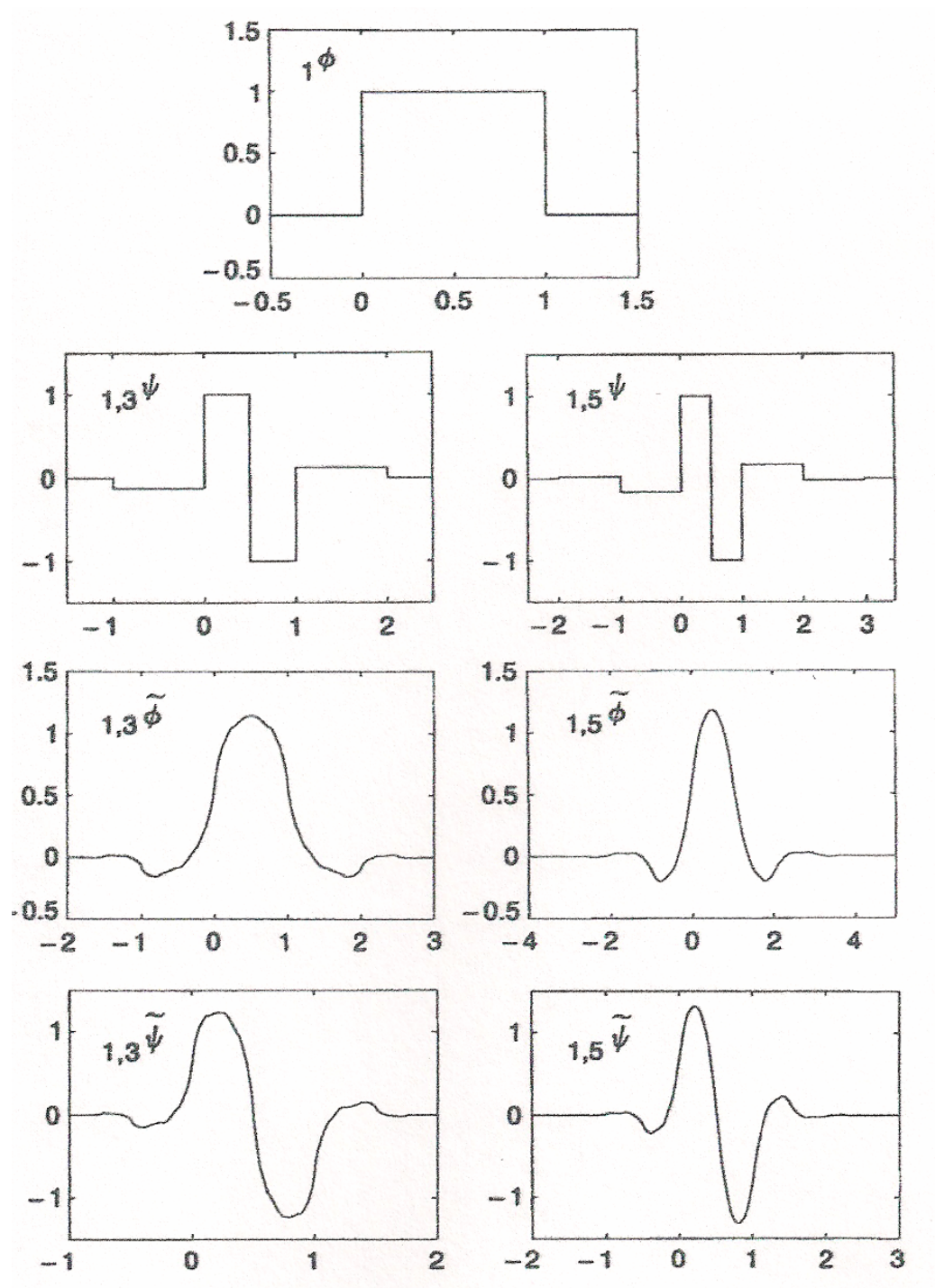
$${}_{N,\tilde{N}}m_0(\xi) = e^{-ik\xi/2} (\cos \xi/2)^{\tilde{N}} \left[\sum_{n=0}^{k-1} \binom{k-1+n}{n} (\sin \xi/2)^{2n} + (\sin \xi/2)^{2k} R(\cos \xi) \right],$$

donde $\tilde{N} \geq 1$, $N + \tilde{N} = 2k$, y R es un polinomio par.

La siguiente lista muestra los coeficientes de los polinomios ${}_Nm_0(\xi)$ y ${}_{N,\tilde{N}}\tilde{m}_0(\xi)$ para los órdenes N y $\tilde{N}$ de las funciones $B - spline$ y las respectivas gráficas ${}_N\varphi(x)$, ${}_{N,\tilde{N}}\tilde{\varphi}(x)$, ${}_N\psi(x)$, ${}_{N,\tilde{N}}\tilde{\psi}$. Los correspondientes términos de los filtros ${}_Nh_k$ y ${}_{N,\tilde{N}}\tilde{h}_k$ se obtiene multiplicando por $\sqrt{2}$ los coeficientes de z^k en ${}_Nm_0$ y ${}_{N,\tilde{N}}\tilde{m}_0$ respectivamente.

N	${}_N m_0$	$\tilde{N}$	$_{N,\tilde{N}} \tilde{m}_0$
1	$\frac{1}{2}(1+z)$	1	$\frac{1}{2}(1+z)$
		3	$-\frac{z^{-2}}{16} + \frac{z^{-1}}{16} + \frac{1}{2} + \frac{z}{2} + \frac{z^2}{16} - \frac{z^3}{16}$
		5	$\frac{3z^{-4}}{256} - \frac{3z^{-3}}{256} - \frac{11z^{-2}}{128} + \frac{11z^{-1}}{128} + \frac{1}{2} + \frac{z}{2} + \frac{11z^2}{128} - \frac{11z^3}{128} - \frac{3z^3}{256} - \frac{3z^4}{256} + \frac{3z^5}{256}$
3	$\frac{1}{8}(z^{-1} + 3 + 3z + z^2)$	1	$-\frac{z^{-1}}{4} + \frac{3}{4} + \frac{3z}{4} - \frac{3z^4}{4}$
		3	$\frac{3z^{-3}}{64} - \frac{9z^{-2}}{64} - \frac{7z^{-1}}{64} + \frac{45}{64} + \frac{45z}{64} - \frac{7z^2}{64} - \frac{9z^3}{64} + \frac{3z^4}{64}$
		7	$2^{-14}(35z^{-7} - 105z^{-6} - 195z^{-5} + 865z^{-4} + 336z^{-3} - 3489z^{-2} - 307z^{-1} + 11025 + 11025z \dots)$

Los correspondientes términos de los filtros ${}_N h_k$ y $_{N,\tilde{N}} \tilde{h}_k$ se obtiene multiplicando $\sqrt{2}$ los coeficientes de z^k en ${}_N m_0$ y $_{N,\tilde{N}} \tilde{m}_0$ respectivamente.



BIORTHOGONAL BASES

54

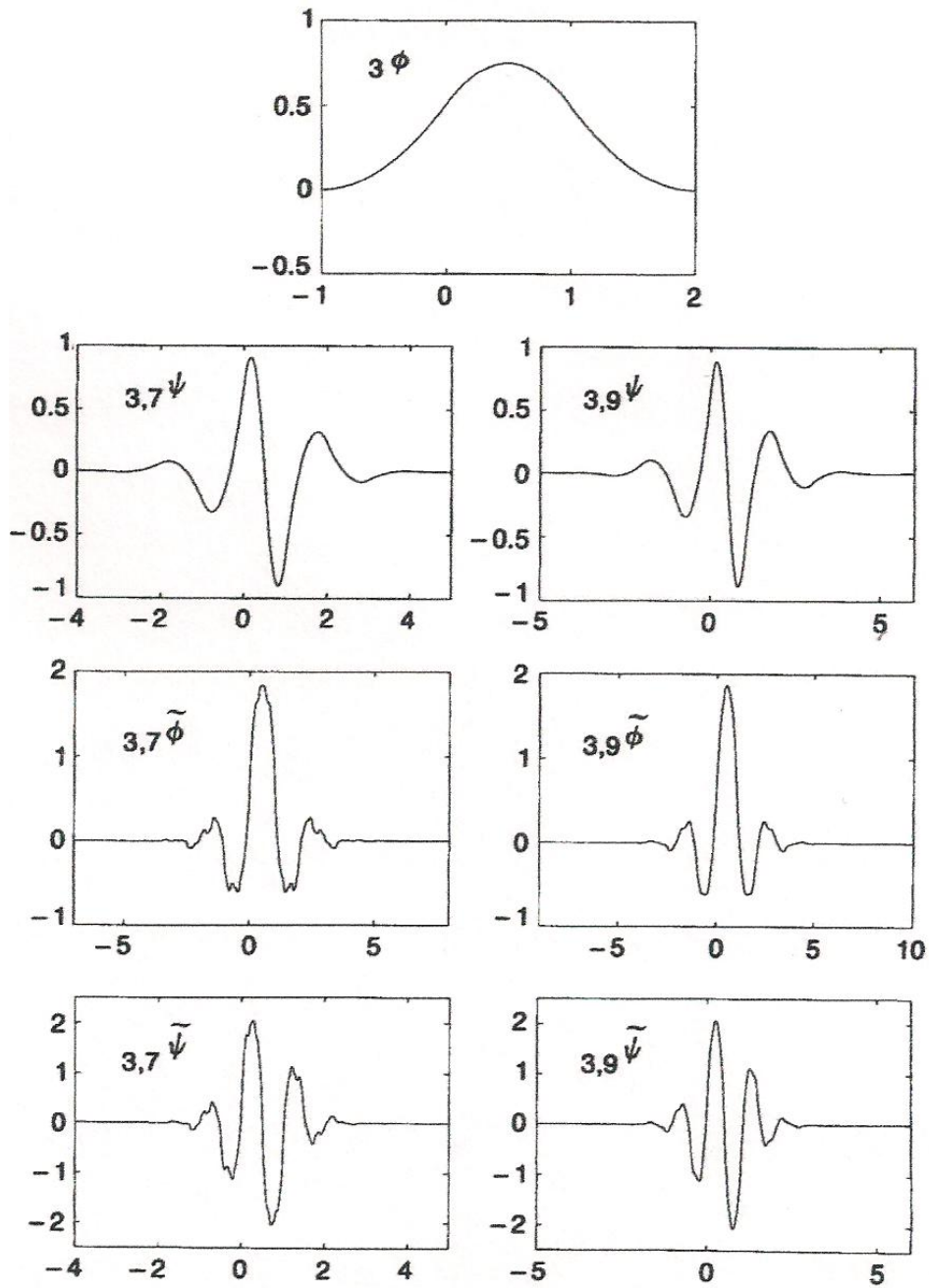


Figure 6.3 (continued).

CAPÍTULO 3

La ecuación de Korteweg-de Vries (KdV)

3.1. Introducción

El objetivo de este capítulo es presentar algunos elementos y propiedades de la ecuación de Korteweg-de Vries (KdV)

$$\frac{\partial u}{\partial t} + 6u \frac{\partial u}{\partial x} + \frac{\partial^3 u}{\partial x^3} = 0, \quad x \in \mathbb{R}, \quad t > 0 \quad (3.1.1)$$

el número 6 es sólo un factor de escala que permite que las soluciones sean fáciles de describir, esta ecuación aparece en el estudio de ondas en aguas poco profundas en la dinámica de fluidos [27, 44, 52, 61]. Una propiedad que caracteriza la ecuación KdV es que los términos uu_x , no linealidad, y u_{xxx} , la dispersión, se balancean entre sí generando soluciones de onda que se propagan sin cambiar de forma.

La ecuación KdV se dedujo en 1895 por D. J. Korteweg y G. de Vries para modelar las ondas de agua en un canal de poca profundidad, con el propósito de resolver una conjetura formulada por el ingeniero John Scott Russell en 1834, donde él afirmaba que en este tipo de canal se genera una ola solitaria, pero varios matemáticos destacados, incluyendo Stokes, estaban convencidos de que tal fenómeno era imposible. Korteweg y de Vries demostraron que Russell estaba en lo cierto, encontrando en forma explícita

y cerrada soluciones de onda viajera a su ecuación con características muy especiales como retención de su forma en todo momento, bien localizada (asintóticamente constante en $\pm\infty$) y en el caso que una onda solitaria traspase a otra, se retiene el tamaño y su forma. Este hecho es en cierta forma el principio de superposición.

La ecuación KdV no recibió mayor atención hasta 1965, cuando N. Zabusky y M. Kruskal publicaron sus resultados de los experimentos numéricos con esta ecuación [71]. Ellos obtuvieron soluciones aproximadas a la ecuación KdV mostrando que cualquier perfil de onda inicial localizada que evoluciona de acuerdo a la ecuación KdV, genera un conjunto finito de ondas viajeras localizadas de la misma forma que la onda solitaria original, confirmando lo descubierto por Korteweg y de Vries en 1895. El término solitón fue acuñado por Zabusky y Kruskal para describir esta onda solitaria solución de la ecuación KdV, en analogía con otras partículas elementales tales como electrón, protón, fotón, etc. Así, por solitón entendemos una onda solitaria en forma de un pulso que es capaz de trasladarse sin cambio de forma y sin pérdidas de energía, y además es capaz de conservar su estructura después de un choque con su semejante, es decir, su comportamiento es análogo al de una partícula. Luego los solitones son ondas no lineales que exhiben un comportamiento extremadamente inesperado e interesante, ondas solitarias que se propagan sin deformarse [14, 28, 52].

3.2. Ondas dispersivas

Recordemos algunas propiedades básicas sobre ondas lineales, tales como dispersión, número de onda y frecuencia, relación fundamental para mostrar la importancia de distinguir entre la velocidad de fase y la velocidad de grupo. Para ello, consideremos la ecuación de onda

$$\frac{\partial^2 u}{\partial t^2} = c^2 \Delta u = c^2 \sum_{j=1}^n \frac{\partial^2 u}{\partial x_j^2},$$

donde c es una constante que representa la velocidad de propagación y Δu es el laplaciano de u . En el caso unidimensional tenemos la ecuación de la cuerda vibrante

$$\frac{\partial^2 u}{\partial t^2} = c^2 \frac{\partial^2 u}{\partial x^2} \tag{3.2.1}$$

cuya solución general es dada por

$$u(x, t) = f(x - ct) + g(x + ct),$$

con f y g funciones arbitrarias, véase p.e., [30], [44] o [61]. Físicamente, f y g representan ondas que se mueven con velocidad constante c y sin cambio de forma a lo largo del eje x , tanto en la parte positiva como en la negativa.

Las soluciones f y g corresponden a los dos factores cuando la ecuación unidimensional (3.2.1) se escribe en la forma

$$\left(\frac{\partial}{\partial t} + c\frac{\partial}{\partial x}\right)\left(\frac{\partial}{\partial t} - c\frac{\partial}{\partial x}\right)u = 0.$$

La ecuación lineal de onda más simple es

$$u_t + cu_x = 0$$

y su solución $u(x, t) = f(x - ct)$ representa una onda que se mueve en la dirección del eje x positivo, con velocidad constante c y sin cambio de forma.

La onda lineal más elemental solución de la ecuación (3.2.1) es la plana o armónica, cuya forma es

$$u(x, t) = Ae^{i(kx - \omega t)}, \quad (3.2.2)$$

donde A , k y ω son constantes; k es el número de onda y se relaciona con la longitud de onda, λ , por $k = \frac{2\pi}{\lambda} = \frac{\omega}{c}$, ω es la frecuencia angular y A es la amplitud. Claramente la frecuencia angular y el número onda tienen la relación de dispersión $\omega = ck$. En general, la relación de dispersión es una expresión de la forma $\omega = \omega(k)$. En la propagación unidimensional de ondas, tenemos

$$e^{i(kx - \omega t)} = e^{ik(x - \frac{\omega}{k}t)},$$

que representa una onda viajera con velocidad de onda $= \frac{\omega(k)}{k}$. Si la velocidad depende realmente del número de onda, entonces ondas con diferente longitud se mueven con distinta velocidad. Luego un problema de propagación de ondas es dispersivo si la velocidad de onda depende del número de onda y, en particular, si la velocidad de onda no es constante. En consecuencia, un problema es no dispersivo, si $\frac{\omega(k)}{k} = \text{constante}$, o lo que es lo mismo, $\omega(k) = \text{constante} \times k$. Otros términos importantes en esta teoría son los de velocidad de fase y velocidad de grupo, esto es, la velocidad de fase de una

50 La ecuación de Korteweg-de Vries (KdV)

onda de la forma $u(x, t) = f(kx - \omega t)$, con k y ω constantes, se define por la constante

$$c_p := \frac{\omega}{k}$$

y la velocidad de grupo de un bloque o paquete de ondas se define por la relación

$$c_g := \frac{d\omega}{dk}.$$

Nótese que para el caso que estamos tratando, se tiene

$$c_p = \frac{\omega}{k} = \frac{ck}{k} = c = \frac{d\omega}{dk} = c_g,$$

es decir, tenemos una relación no dispersiva. Así, un problema unidimensional de onda es dispersivo si $\frac{d\omega}{dk} \neq \text{constante}$, véase p.e., [8], [27], [28] o [44].

Para la ecuación KdV linealizada

$$\frac{\partial u}{\partial t} + c \frac{\partial u}{\partial x} + \beta \frac{\partial^3 u}{\partial x^3} = 0$$

donde se presenta una mayor dispersión, debido a la gran longitud de onda en comparación con el resto de longitudes características del modelo físico. La relación de dispersión para la onda viajera

$$u(x, t) = e^{i(kx - \omega t)}$$

está dada por

$$\omega = ck - \beta k^3,$$

nótese que esta relación de dispersión es no lineal entre la frecuencia ω y el número de onda k . De las expresiones para las velocidades de fase y de grupo tenemos

$$c_p = \frac{\omega}{k} = c - \beta k^2 \quad \text{y} \quad \frac{d\omega}{dk} = c - 3\beta k^2 = c_g,$$

de lo cual se deduce que $\frac{d\omega}{dk} \neq \text{constante}$; de donde la ecuación KdV linealizada, es de onda dispersiva. Observe también que si $\beta > 0$ se tiene $c_g \leq c_p$.

3.3. Derivación de la ecuación KdV

La ecuación de Korteweg-de Vries (3.1.1) admite solución tipo onda solitaria o solitón. Existen varios métodos para obtener tal solución, presentaremos algunos de ellos en este trabajo por motivos de completitud.

Comencemos presentando la solución canónica de la ecuación KdV, para ello supongamos una solución de onda viajera con la estructura

$$u(x, t) = v(x - ct), \quad x \in \mathbb{R}, \quad t > 0,$$

para alguna función v y velocidad de onda constante c . Determinemos v por sustitución de esta expresión en la ecuación (3.1.1).

Sea $\xi = x - ct$, entonces u es solución de la ecuación KdV siempre que v satisfaga la ecuación diferencial ordinaria

$$-c \frac{dv}{d\xi} + 6v \frac{dv}{d\xi} + \frac{d^3 v}{d\xi^3} = 0$$

al integrar se obtiene

$$-cv + 3v^2 + \frac{d^2 v}{d\xi^2} = A,$$

donde A es una constante de integración. Al multiplicar esta última expresión por $\frac{dv}{d\xi}$

$$-cv \frac{dv}{d\xi} + 3v^2 \frac{dv}{d\xi} + \frac{d^2 v}{d\xi^2} \frac{dv}{d\xi} = A \frac{dv}{d\xi},$$

e integrar, se tiene

$$-\frac{c}{2}v^2 + v^3 + \frac{1}{2} \left(\frac{dv}{d\xi} \right)^2 = Av + B,$$

con B una constante de integración. Como la onda solución es un solitón, entonces cuando $|\xi| \rightarrow \infty$, $v, \frac{dv}{d\xi}, \frac{d^2 v}{d\xi^2} \rightarrow 0$, de donde $A = B = 0$. Luego

$$-\frac{c}{2}v^2 + v^3 + \frac{1}{2} \left(\frac{dv}{d\xi} \right)^2 = 0 \quad \text{o} \quad \frac{dv}{d\xi} = \pm v \sqrt{c - 2v}.$$

Al separar variables e integrar se obtiene

$$\xi = - \int \frac{dv}{v \sqrt{c - 2v}} + d,$$

52 La ecuación de Korteweg-de Vries (KdV)

d una constante de integración. Haciendo la sustitución

$$v = \frac{c}{2} \operatorname{sech}^2 y \quad dv = -c \operatorname{sech}^2 y \tanh y dy$$

y

$$v\sqrt{c-2v} = \frac{c^{3/2}}{2} \operatorname{sech}^2 y \tanh y$$

se obtiene

$$-\int \frac{dv}{v\sqrt{c-2v}} = \frac{2}{\sqrt{c}} y + d.$$

Por tanto,

$$\xi = \frac{2}{\sqrt{c}} y + d \quad \text{o} \quad y = \frac{\sqrt{c}}{2} \xi - d,$$

pero $y = \operatorname{sech}^{-1} \sqrt{\frac{2v}{c}}$, luego

$$\operatorname{sech}\left(\frac{\sqrt{c}}{2} \xi - d\right) = \sqrt{\frac{2v}{c}}$$

de donde

$$v(\xi) = \frac{c}{2} \operatorname{sech}^2\left(\frac{\sqrt{c}}{2} \xi - d\right).$$

En consecuencia, la solución de la ecuación KdV es

$$u(x, t) = \frac{c}{2} \operatorname{sech}^2\left(\frac{\sqrt{c}}{2} (x - ct) - d\right).$$

Como la ecuación KdV es completamente integrable [33, 51], lo que significa que la solución exacta se puede calcular para un valor inicial arbitrario, y además, por la naturaleza de la solución, se puede elegir $d = 0$. Obteniéndose así,

$$u(x, t) = \frac{c}{2} \operatorname{sech}^2\left(\frac{\sqrt{c}}{2} (x - ct)\right).$$

En virtud a que $v > 0$ para todo ξ , el solitón es una onda de elevación simétrica sobre $\xi = 0$, que se propaga en el medio sin cambio de forma con velocidad constante c proporcional a la amplitud. Por lo tanto, ondas solitarias con gran amplitud se mueven con mayor velocidad que ondas solitarias con menor amplitud.

Si en lugar de $v = \frac{c}{2} \operatorname{sech}^2 y$ hacemos $v = -\frac{c}{2} \operatorname{csch}^2 y$, obtenemos otra solución de la ecuación KdV del tipo onda viajera dada por

$$u(x, t) = -\frac{c}{2} \operatorname{csch}^2 \left(\frac{\sqrt{c}}{2} (x - ct) \right).$$

Nótese que esta solución no es un solitón, debido a que $u(x, t)$ no es acotada en $\xi = 0$.

La ecuación KdV también se puede resolver utilizando la transformación de Bäcklund, para más detalles véase por ejemplo, [27], [28] o [52]. Esta transformación consiste en introducir una función v tal que $u = v_x$, donde $v_x = \frac{\partial v}{\partial x}$. De la ecuación (3.1.1) se sigue

$$v_{tx} + 6v_x v_{xx} + v_{xxx} = \frac{\partial}{\partial x} [v_t + 3v_x^2 + v_{xxx}] = 0,$$

al integrar con respecto a x se tiene

$$v_t + 3v_x^2 + v_{xxx} = g(t), \quad (3.3.1)$$

donde $g(t)$ es una función de integración. El cambio de variable

$$v^* = v - \int^t g(r) dr$$

elimina g de la ecuación (3.3.1). Como la solución u de (3.1.1) se obtendrá de la solución v de (3.3.1) y $v_x^* = v_x$, entonces no necesitamos hacer distinción entre v^* y v . Luego podemos hacer $g = 0$ sin pérdida de generalidad.

Existe una transformación de Bäcklund que deja invariante la ecuación. A este tipo de transformación se le llama auto-transformación de Bäcklund, que para la ecuación de Korteweg-de Vries está dada por el conjunto de ecuaciones

$$\begin{aligned} \tilde{v}_\xi &= -v_x + \beta - \frac{1}{2}(v - \tilde{v})^2 \\ \tilde{v}_\tau &= -v_t + (v - \tilde{v})(v_{xx} - \tilde{v}_{\xi\xi}) - 2(v_x^2 + v_x \tilde{v}_\xi + \tilde{v}_\xi^2) \\ \xi &= x, \quad \tau = t, \end{aligned} \quad (3.3.2)$$

donde β es un parámetro. Ahora generamos una solución no trivial de la ecuación KdV aplicando la transformación (3.3.2) con $\tilde{v} = 0$

$$\begin{aligned} v_x &= \beta - \frac{1}{2}v^2 \\ v_t &= vv_{xx} - 2v_x^2. \end{aligned} \quad (3.3.3)$$

54 La ecuación de Korteweg-de Vries (KdV)

La solución de la primera ecuación se obtiene haciendo $v(x, t) = z(s)$ con $s = x - ct$. Así,

$$v_x = \frac{dz}{ds} = \beta - \frac{1}{2}z^2$$

al separar variables e integrar se tiene

$$\int \frac{dz}{\beta - \frac{1}{2}z^2} = \int ds.$$

Al hacer la sustitución

$$z = \sqrt{2\beta} \tanh y, \quad dz = \sqrt{2\beta} \operatorname{sech}^2 y dy$$

$$\beta - \frac{1}{2}z^2 = \beta(1 - \tanh^2 y) = \beta \operatorname{sech}^2 y,$$

luego

$$\int \frac{dz}{\beta - \frac{1}{2}z^2} = \frac{\sqrt{2\beta}}{\beta} \int dy = \sqrt{\frac{2}{\beta}} y.$$

Por tanto,

$$\sqrt{\frac{2}{\beta}} y = s + c_1$$

c_1 es una constante de integración, la cual se puede hacer cero sin pérdida de generalidad. Así,

$$y = \sqrt{\frac{\beta}{2}} s = \tanh^{-1} \left(\frac{z}{\sqrt{2\beta}} \right)$$

o lo que es lo mismo,

$$z = \sqrt{2\beta} \tanh \left(\sqrt{\frac{\beta}{2}} s \right),$$

luego

$$v(x, t) = \sqrt{2\beta} \tanh \left(\sqrt{\frac{\beta}{2}} (x - ct) \right)$$

pero $u = v_x$, por tanto

$$u(x, t) = \beta \operatorname{sech}^2 \left(\sqrt{\frac{\beta}{2}} (x - ct) \right)$$

al hacer $\beta = c/2$ se tiene

$$u(x, t) = \frac{c}{2} \operatorname{sech}^2\left(\frac{\sqrt{c}}{2} (x - ct)\right).$$

Con un procedimiento igual se obtiene

$$v(x, t) = \sqrt{2\beta} \coth\left(\sqrt{\frac{\beta}{2}} (x - ct)\right)$$

de la que se deduce

$$u(x, t) = -\frac{c}{2} \operatorname{csch}^2\left(\frac{\sqrt{c}}{2} (x - ct)\right).$$

3.4. El método de Lax

En sus trabajos pioneros sobre sistemas completamente integrables, Gardner et al. [33] formularon el método de la transformada inversa scattering para encontrar soluciones exactas de la ecuación KdV. Lax utiliza este método con ligeras modificaciones para encontrar soluciones de ecuaciones de evolución no lineales con todas sus integrales [45]. Su formulación tiene la característica de asociación de ciertas ecuaciones de evolución no lineales con ecuaciones lineales, como son los análogos de la ecuación de Schrödinger para la ecuación KdV.

En la formulación del método de Lax, se consideran dos operadores lineales L y M . El primer operador define un problema espectral para la función auxiliar v , en la variable espacial x y con valor espectral λ ; mientras que la variable t se considera como un parámetro. En otras palabras, la ecuación característica asociada con el operador L es entonces

$$Lv = \lambda v. \quad (3.4.1)$$

El operador M describe el cambio de los valores propios con el parámetro t , que usualmente representa el tiempo en una ecuación de evolución no lineal. La forma general de esta ecuación de evolución es

$$v_t = Mv. \quad (3.4.2)$$

Derivando la ecuación característica (3.4.1) con respecto a t se tiene

$$L_t v + Lv_t = \lambda_t v + \lambda v_t. \quad (3.4.3)$$

56 La ecuación de Korteweg-de Vries (KdV)

Al reemplazar la ecuación (3.4.2) en (3.4.3), eliminamos v_t y obtenemos

$$L_t v + L M v = \lambda_t v + \lambda M v = \lambda_t v + M \lambda v = \lambda_t v + M L v$$

o equivalentemente,

$$\lambda_t v = (L_t + L M - M L) v. \quad (3.4.4)$$

Por tanto, los valores propios son constantes ($\lambda_t = 0$) para las funciones propias no nulas, si y sólo si

$$\frac{\partial L}{\partial t} = M L - L M = -[L, M], \quad (3.4.5)$$

donde $[L, M] = L M - M L$ es el conmutador de los operadores L y M . L_t se interpreta como la derivada del operador, y no la derivada del operador actuando sobre una función. Los operadores L y M , y la ecuación (3.4.5), se llaman par de Lax y ecuación de Lax, respectivamente.

En general, no existe un método sistemático para determinar los operadores L y M para una ecuación de evolución dada. Sin embargo, para la ecuación KdV estos operadores están dados por

$$L = -\frac{\partial^2}{\partial x^2} - u \quad (3.4.6)$$

$$M = 4\frac{\partial^3}{\partial x^3} + 6u\frac{\partial}{\partial x} + 3\frac{\partial u}{\partial x},$$

donde $u = u(x, t)$ [27]. Veamos que se cumple la ecuación (3.4.5).

$$\begin{aligned} L M &= -\left[\frac{\partial^2}{\partial x^2} + u\right] \left[4\frac{\partial^3}{\partial x^3} + 6u\frac{\partial}{\partial x} + 3\frac{\partial u}{\partial x}\right] \\ &= -4\frac{\partial^5}{\partial x^5} - 6\left(\frac{\partial^2 u}{\partial x^2} \frac{\partial}{\partial x} + 2\frac{\partial u}{\partial x} \frac{\partial^2}{\partial x^2} + u\frac{\partial^3}{\partial x^3}\right) \\ &\quad - 3\left(\frac{\partial^3 u}{\partial x^3} + 2\frac{\partial^2 u}{\partial x^2} \frac{\partial}{\partial x} + \frac{\partial u}{\partial x} \frac{\partial^2}{\partial x^2}\right) - 4u\frac{\partial^3}{\partial x^3} - 6u^2\frac{\partial}{\partial x} - 3u\frac{\partial u}{\partial x} \end{aligned}$$

$$\begin{aligned} M L &= -\left[4\frac{\partial^3}{\partial x^3} + 6u\frac{\partial}{\partial x} + 3\frac{\partial u}{\partial x}\right] \left[\frac{\partial^2}{\partial x^2} + u\right] \\ &= -4\frac{\partial^5}{\partial x^5} - 4\left(\frac{\partial^3 u}{\partial x^3} + 3\frac{\partial^2 u}{\partial x^2} \frac{\partial}{\partial x} + 3\frac{\partial u}{\partial x} \frac{\partial^2}{\partial x^2} + u\frac{\partial^3}{\partial x^3}\right) \\ &\quad - 6u\left(\frac{\partial^3}{\partial x^3} + \frac{\partial u}{\partial x} + u\frac{\partial}{\partial x}\right) - 3\frac{\partial u}{\partial x} \frac{\partial^2}{\partial x^2} - 3u\frac{\partial u}{\partial x}. \end{aligned}$$

Por lo tanto,

$$L_t = ML - LM = -\frac{\partial^3 u}{\partial x^3} - 6u \frac{\partial u}{\partial x}$$

Pero $\frac{\partial L}{\partial t} = u_t$, luego

$$u_t + 6uu_x + u_{xxx} = 0.$$

Terminemos esta sección comentando que la ecuación KdV (3.1.1) pertenece a una clase de ecuaciones de onda no lineal, que tiene la propiedad del par de Lax, como también la integrabilidad a través de la transformada inversa scattering, véase por ejemplo [27], [28], [33] o [50]. La propiedad de la integrabilidad se caracteriza por la existencia de un número infinito de leyes de conservación independientes tales como

$$\begin{aligned} \int_{-\infty}^{\infty} u dx &= \text{constante} \\ \int_{-\infty}^{\infty} u^2 dx &= \text{constante} \\ \int_{-\infty}^{\infty} \left(u^3 - \frac{1}{2}u_x^2\right) dx &= \text{constante} \end{aligned}$$

que se pueden asociar con la conservación de masa, momentum y energía, respectivamente. La primera de estas expresiones se obtiene integrando la ecuación KdV (3.1.1) sobre todo el eje x y teniendo en cuenta que u y sus derivadas tienden a cero cuando $|x| \rightarrow \infty$. La segunda ley de conservación se obtiene multiplicando la ecuación KdV por u , y escribirla en la forma

$$\frac{\partial}{\partial t} \left(\frac{1}{2}u^2 \right) + \frac{\partial}{\partial x} \left(uu_{xx} - \frac{1}{2}u_x^2 + 2u^3 \right) = 0.$$

De donde

$$\int_{-\infty}^{\infty} \frac{1}{2}u^2 dx = \text{constante}$$

para todas las soluciones de la ecuación (3.1.1) que se anulan lo suficientemente rápido en el infinito. La tercera ley de conservación se obtiene multiplicando la ecuación KdV por $3u^2$ y luego restar u_x por $\frac{\partial}{\partial x}(\text{KdV})$, es decir,

$$\begin{aligned} 3u^2(\text{KdV}) - u_x \times \frac{\partial}{\partial x}(\text{KdV}) &= \frac{\partial}{\partial t} \left(u^3 - \frac{1}{2}u_x^2 \right) \\ - \frac{\partial}{\partial x} \left[-\frac{9}{2}u^4 - 3u^2u_{xx} + 6uu_x^2 - u_xu_{xx} + \frac{1}{2}u_{xx}^2 \right] &= 0. \end{aligned}$$

Así tenemos la tercera constante de movimiento,

$$\int_{-\infty}^{\infty} \left(u^3 - \frac{1}{2} u_x^2 \right) dx = \text{constante}.$$

3.5. Método de Adomian para la ecuación KdV

El método de descomposición de Adomian es una técnica que da soluciones en forma de serie con rápida convergencia a diferentes tipos de ecuaciones que modelan una variedad de problemas en física, biología y química. Este método lo introdujo Adomian en la década de 1980, en la solución de ciertos problemas de física matemática [1, 2]. Es de resaltar que el método se aplica a ecuaciones en derivadas parciales no lineales, como la ecuación de Korteweg-de Vries [39], en donde ésta se descompone en un término lineal y otro no lineal. Para precisar consideremos la ecuación KdV

$$u_t + \mu u u_x + \epsilon u_{xxx} = 0 \quad (3.5.1)$$

con la condición

$$u(x, 0) = u_0(x) = 6x,$$

acá μ y ϵ son constantes positivas.

Definamos el operador lineal $L_t := \frac{\partial}{\partial t}$, entonces la ecuación (3.5.1) se puede expresar como

$$L_t(u) + \mu N(u) + \epsilon u_{xxx} = 0,$$

donde $N(u) = uu_x$ es el término no lineal. Sea

$$L_t^{-1}(\cdot) = \int_0^t (\cdot) d\tau$$

el operador inverso de L_t , entonces aplicando formalmente el operador L_t^{-1} se tiene

$$L_t^{-1} L_t(u) = \int_0^t \frac{\partial u(x, \tau)}{\partial \tau} d\tau = u(x, t) - u(x, 0)$$

de donde

$$u(x, t) = u_0(x) - \mu L_t^{-1} N(u) + \epsilon L_t^{-1} u_{xxx}. \quad (3.5.2)$$

El método de descomposición de Adomian consiste en suponer una solución en forma de serie

$$u(x, t) = \sum_{m=0}^{\infty} u_m(x, t)$$

donde los términos $u_m(x, t)$ son determinados por recurrencia. El término no lineal $N(u)$ se descompone en una serie infinita de polinomios de la forma

$$N(u) = \sum_{m=0}^{\infty} A_m$$

donde A_m son los polinomios de Adomian.

Al sustituir esta descomposición en (3.5.2) se tiene

$$\sum_{m=0}^{\infty} u_m(x, t) = u_0(x) - \mu \sum_{m=0}^{\infty} L_t^{-1} A_m - \epsilon \sum_{m=0}^{\infty} L_t^{-1} \frac{\partial^3 u_m(x, t)}{\partial x^3}.$$

Las soluciones se calculan a través de las fórmulas de recurrencia de la siguiente manera

$$u(x, 0) = u_0(x) = 6x$$

$$u_{k+1}(x, t) = -\mu L_t^{-1} A_k - \epsilon L_t^{-1} \frac{\partial^3 u_k(x, t)}{\partial x^3}, \quad k = 0, 1, 2, \dots$$

Los polinomios de Adomian se definen por la expresión

$$A_m = \frac{1}{m!} \frac{d^m}{d\lambda^m} \left[N \left(\sum_{p=0}^{\infty} \lambda^p u_p \right) \right]_{\lambda=0}, \quad \text{para } m = 0, 1, 2, \dots$$

De esta manera, los polinomios de Adomian están dados por

$$A_0 = N(u_0) = u_0 \frac{\partial u_0}{\partial x} = u_0 u_{0_x} = (6x)(6) = 36x.$$

Ahora calculemos los demás polinomios de Adomian

$$\begin{aligned} A_1 &= \frac{d}{d\lambda} \left[N(u_0 + \lambda u_1) \right]_{\lambda=0} \\ &= \frac{d}{d\lambda} \left[(u_0 + \lambda u_1)(u_0 + \lambda u_1)_x \right]_{\lambda=0} \\ &= \frac{d}{d\lambda} (u_0 + \lambda u_1)(u_{0_x} + \lambda u_{1_x}) \Big|_{\lambda=0} \\ &= u_1(u_{0_x} + \lambda u_{1_x}) + u_{1_x}(u_0 + \lambda u_1) \Big|_{\lambda=0} \\ &= u_1 u_{0_x} + u_0 u_{1_x} \end{aligned}$$

$$\begin{aligned}
 A_2 &= \frac{1}{2} \frac{d^2}{d\lambda^2} \left[N(u_0 + \lambda u_1 + \lambda^2 u_2) \right]_{\lambda=0} \\
 &= \frac{1}{2} \frac{d^2}{d\lambda^2} (u_0 + \lambda u_1 + \lambda^2 u_2)(u_{0x} + \lambda u_{1x} + \lambda^2 u_{2x}) \Big|_{\lambda=0} \\
 &= u_0 u_{2x} + u_1 u_{1x} + u_2 u_{0x}
 \end{aligned}$$

$$\begin{aligned}
 A_3 &= \frac{1}{3!} \frac{d^3}{d\lambda^3} \left[N(u_0 + \lambda u_1 + \lambda^2 u_2 + \lambda^3 u_3) \right]_{\lambda=0} \\
 &= u_0 u_{3x} + u_1 u_{2x} + u_2 u_{1x} + u_3 u_{0x}
 \end{aligned}$$

$$\begin{aligned}
 A_4 &= \frac{1}{4!} \frac{d^4}{d\lambda^4} \left[N(u_0 + \lambda u_1 + \lambda^2 u_2 + \lambda^3 u_3 + \lambda^4 u_4) \right]_{\lambda=0} \\
 &= u_0 u_{4x} + u_1 u_{3x} + u_2 u_{2x} + u_3 u_{1x} + u_4 u_{0x}
 \end{aligned}$$

$$\begin{aligned}
 A_5 &= \frac{1}{5!} \frac{d^5}{d\lambda^5} \left[N(u_0 + \lambda u_1 + \lambda^2 u_2 + \lambda^3 u_3 + \lambda^4 u_4 + \lambda^5 u_5) \right]_{\lambda=0} \\
 &= u_0 u_{5x} + u_1 u_{4x} + u_2 u_{3x} + u_3 u_{2x} + u_4 u_{1x} + u_5 u_{0x}
 \end{aligned}$$

Encontremos algunos términos de la serie solución

$$u_0(x, t) = u_0(x) = 6x$$

$$\begin{aligned}
 u_1(x, t) &= -\mu \int_0^t A_0 d\tau - \epsilon \int_0^t \frac{\partial^3 u_0(x, \tau)}{\partial x^3} d\tau \\
 &= -\mu t u_0 u_{0x} - \epsilon t \frac{\partial^3 u_0(x, t)}{\partial x^3} = 36\mu xt.
 \end{aligned}$$

Nótese que

$$\frac{\partial^3 u_0(x, t)}{\partial x^3} = 0 \quad \text{y} \quad \frac{\partial^3 u_1(x, t)}{\partial x^3} = 0.$$

$$\begin{aligned}
 u_2(x, t) &= -\mu \int_0^t A_1 d\tau - \epsilon \int_0^t \frac{\partial^3 u_1(x, \tau)}{\partial x^3} d\tau \\
 &= -\mu \int_0^t [(6^2 \mu x \tau)(6) + (6x)(-6^2 \mu \tau)] d\tau \\
 &= 6^3 \mu^2 x \int_0^t 2 \tau d\tau = 6^3 \mu^2 x t^2.
 \end{aligned}$$

Ahora bien,

$$\frac{\partial u_2(x, t)}{\partial x} = 6^3 \mu^2 t^2 \quad \text{y} \quad \frac{\partial^n u_2(x, t)}{\partial x^n} = 0, \quad \text{para } n = 2, 3, 4, \dots$$

$$\begin{aligned} u_3(x, t) &= -\mu \int_0^t A_2 d\tau - \epsilon \int_0^t \frac{\partial^3 u_2}{\partial x^3} d\tau \\ &= -\mu \int_0^t [u_0 u_{2_x} + u_1 u_{1_x} + u_2 u_{0_x}] d\tau - \epsilon \int_0^t 0 d\tau \\ &= -\mu \int_0^t [u_0 u_{2_x} + u_1 u_{1_x} + u_2 u_{0_x}] d\tau \\ &= -\mu \int_0^t [(6x)(6^3 \mu^2 \tau^2) + (-6^2 \mu x \tau)(-6^2 \mu \tau) + (6^3 \mu^2 x \tau^2)(6)] d\tau \\ &= -6^4 \mu^3 x \int_0^t [\tau^2 + \tau^2 + \tau^2] d\tau \\ &= -6^4 \mu^3 x \int_0^t [3\tau^2] d\tau = -6^4 \mu^3 x t^3. \end{aligned}$$

Nuevamente se tiene

$$\frac{\partial u_3(x, t)}{\partial x} = -6^4 \mu^3 t^3 \quad \text{y} \quad \frac{\partial^n u_3(x, t)}{\partial x^n} = 0, \quad \text{para } n = 2, 3, 4, \dots$$

$$\begin{aligned} u_4(x, t) &= -\mu \int_0^t A_3 d\tau - \epsilon \int_0^t \frac{\partial^3 u_3}{\partial x^3} d\tau \\ &= -\mu \int_0^t [u_0 u_{3_x} + u_1 u_{2_x} + u_2 u_{1_x} + u_3 u_{0_x}] d\tau - \epsilon \int_0^t 0 d\tau \\ &= -\mu \int_0^t [(6x)(6^4 \mu^3 \tau^3) + (-6^2 \mu x \tau)(6^3 \mu^2 \tau^2) \\ &\quad + (6^3 \mu^2 x \tau^2)(-6^2 \mu \tau) + (-6^4 \mu^3 x \tau^3)(6)] d\tau \\ &= 6^5 \mu^4 x \int_0^t [\tau^3 + \tau^3 + \tau^3 + \tau^3] d\tau \\ &= 6^5 \mu^4 x \int_0^t [4\tau^3] d\tau = 6^5 \mu^4 x t^4. \end{aligned}$$

62 La ecuación de Korteweg-de Vries (KdV)

En consecuencia,

$$\frac{\partial u_4(x, t)}{\partial x} = 6^5 \mu^4 t^4 \quad \text{y} \quad \frac{\partial^n u_4(x, t)}{\partial x^n} = 0, \quad \text{para } n = 2, 3, 4, \dots$$

$$\begin{aligned} u_5(x, t) &= -\mu \int_0^t A_4 d\tau - \epsilon \int_0^t \frac{\partial^3 u_4}{\partial x^3} d\tau \\ &= -\mu \int_0^t [u_0 u_{4x} + u_1 u_{3x} + u_2 u_{2x} + u_3 u_{1x} + u_4 u_{0x}] d\tau - \epsilon \int_0^t 0 d\tau \\ &= -\mu \int_0^t [(6x)(6^5 \mu^4 \tau^4) + (-6^2 \mu x \tau)(-6^4 \mu^3 \tau^3) + (6^3 \mu^2 x \tau^2)(6^3 \mu^2 \tau^2) \\ &\quad + (-6^4 \mu^3 x \tau^3)(-6^2 \mu \tau) + (6^5 \mu^4 x \tau^4)(6)] d\tau \\ &= -6^6 \mu^5 x \int_0^t [\tau^4 + \tau^4 + \tau^4 + \tau^4 + \tau^4] d\tau \\ &= -6^6 \mu^5 x \int_0^t [5\tau^4] d\tau = -6^6 \mu^5 x t^5. \end{aligned}$$

La serie solución es por tanto

$$\begin{aligned} u(x, t) &= \sum_{m=0}^{\infty} u_m(x, t) = 6x - 6^2 \mu x t + 6^3 \mu^2 x t^2 - 6^4 \mu^3 x t^3 + 6^5 \mu^4 x t^4 - 6^6 \mu^5 x t^5 + \dots \\ &= 6x [1 - 6\mu t + (6\mu t)^2 - (6\mu t)^3 + (6\mu t)^4 - (6\mu t)^5 + \dots] \\ &= 6x \sum_{m=0}^{\infty} (-1)^m (6\mu t)^m = \frac{6x}{1 + 6\mu t}, \quad |6\mu t| < 1. \end{aligned}$$

Nótese que si $\mu = 6$ y $\epsilon = 1$ estamos con la ecuación KdV canónica (3.1.1), y la solución es entonces

$$u(x, t) = \frac{6x}{1 + 36t}, \quad |36t| < 1.$$

CAPÍTULO 4

Método Wavelet-Petrov-Galerkin para la ecuación KdV

4.1. Introducción

Al resolver ecuaciones diferenciales por el método de Galerkin, los espacios funcionales contruidos no sólo deben tener buenas propiedades de aproximación, sino que también deben permitir facilidad y rapidez en los cálculos computacionales. La idea básica detrás del método de Galerkin es extremadamente simple, se parte del conocimiento de una base y consiste en construir una sucesión de espacios funcionales $\{V_m\}$ de dimensión finita generado por $\{w_1, w_2, \dots, w_m\}$ de manera que $V_1 \subset V_2 \subset \dots \subset V$, y

$$\lim_{m \rightarrow \infty} d(v, V_m) = 0,$$

donde $d(v, V_m)$ es la distancia de v a V_m , para todo $v \in V$. Esta sucesión de espacios $\{V_m\}$ permite construir la sucesión $\{v_m\}$ siendo v_m la solución única del problema variacional aproximado: hallar $v_m \in V_m$ tal que,

$$a(v_m, v) = l(v) \quad \text{para todo } v \in V. \quad (4.1.1)$$

Esta es la esencia del método de Galerkin. Acá $a(\cdot, \cdot)$ es una forma bilineal y l es un funcional lineal. Nótese que (4.1.1) no es más que un sistema lineal de m ecuaciones algebraicas con m incógnitas, y por tanto para determinar

v_m basta calcular $c_j \in \mathbb{R}$ tal que

$$\begin{aligned} a\left(\sum_{j=1}^m c_j w_j, w_i\right) &= l(w_i) \\ \sum_{j=1}^m a(w_j, w_i) c_j &= l(w_i), \quad i = 1, 2, \dots, m \end{aligned}$$

cuyo sistema homogéneo asociado tiene como solución única la trivial. Luego existe y es única $c_j \in \mathbb{R}$, como también lo es v_m , donde

$$v_m = \sum_{j=1}^m c_j w_j.$$

La solución aproximada v_m es, en general, diferente de la solución exacta v . Al incrementar la precisión, es natural buscar la solución aproximada en un subespacio más grande V_m . Por tanto, para una sucesión de subespacios $V_1 \subset V_2 \subset \dots \subset V$, calculamos la correspondiente sucesión de soluciones aproximadas $v_m \in V_m$, $m = 1, 2, \dots$. El procedimiento para encontrar esta solución, se conoce como *método de Galerkin* [11, 12, 17, 29, 38, 42, 57].

Mencionemos también que bajo ciertas condiciones sobre la forma bilineal $a(\cdot, \cdot)$, y si además, $V_1 \subset V_2 \subset \dots$ es una sucesión de subespacios de un espacio de Hilbert V con la propiedad

$$\overline{\bigcup_{m \geq 1} V_m} = V,$$

entonces el método de Galerkin converge

$$\|v - v_m\|_H \rightarrow 0 \quad \text{cuando} \quad m \rightarrow \infty.$$

En consecuencia, el problema (4.1.1) se reduce a resolver el sistema de ecuaciones lineales

$$\mathbf{K}^t \mathbf{c} = \mathbf{F},$$

en donde $\mathbf{K} = (K_{ij})$ es la matriz de rigidez, siendo $K_{ij} = a(w_i, w_j)$, el vector $l(w_i)$ se llama vector de carga y el vector de componentes c_j es el vector de desplazamientos nodales.

En la solución numérica de problemas de valor de frontera cuando se utilizan las wavelets como funciones admisibles y de prueba en la

aproximación de Galerkin, el esquema resultante se llama *Método Wavelet-Galerkin* (MWG). La ventaja de este método es el preconditionamiento y el potencial de discretización de las wavelets para los algoritmos adaptativos [15]. En los métodos de elementos finitos o diferencias finitas para la solución de problemas elípticos en dominios acotados, una vez discretizado el problema, el sistema $Ax = b$ a resolver, es mal condicionado [17, 29, 57]. Este mal condicionamiento tiene dos inconvenientes: la inestabilidad numérica y la lenta convergencia para algoritmos de resolución iterativos. Con el fin de obviar este problema, usualmente se utiliza un preconditionamiento que consiste en encontrar una matriz invertible D tal que $D^{-1}AD$ con la que se tendrá un mejor número de condición [62]. Jaffard demuestra en [41] que al utilizar métodos wavelets, el preconditionamiento más simple se tiene cuando D es una matriz diagonal, obteniéndose así, $\kappa(A) = O(1)$, donde $\kappa(A)$ es el número de condición. Por tanto, las técnicas wavelets proporcionan métodos numéricos eficientes, como alternativa a los métodos clásicos [7, 22, 23, 34, 53].

En sus trabajos sobre análisis wavelets, Daubechies presenta un método para construir wavelets de soporte compacto y funciones de escala con regularidad arbitraria y momento cero [24, 25]. Sin embargo, el precio para estas buenas propiedades es la carencia de simetría y soporte amplio. Este inconveniente desaparece en el contexto de las wavelets biortogonales, un concepto introducido por Cohen, Daubechies y Feauveau en [19]. Por lo tanto, en este contexto dos funciones base no ortogonales $\psi_{j,k}$ y $\psi_{j,k}^*$ que también se llaman wavelets, son construidas a partir de las funciones de escala trasladadas ψ y ψ^* .

A diferencia del método de Galerkin, en donde las mismas funciones base son usadas tanto como de prueba como admisibles, en el método de Petrov-Galerkin, las funciones de prueba y ensayo son diferentes. En la aproximación de Petrov-Galerkin por wavelets biortogonales, la idea es usar una de las familias de funciones base como admisibles y su dual como funciones de prueba. Dahlke y Weinreich en [20] construyen bases wavelets biortogonales de soporte compacto adaptados a algunos operadores diferenciales para obtener matrices de rigidez con estructuras simples.

El propósito de este capítulo es estudiar la precisión del método de Petrov-Galerkin utilizando wavelets biortogonales. Para tal fin, se estudia la ecuación KdV

$$u_t + \mu uu_x + \epsilon u_{xxx} = 0$$

con la condición inicial $u(x, 0) = u_0(x)$, donde μ y ϵ son constantes positivas. Y en lugar de las bases wavelets multi-nivel, se considera expandir las soluciones aproximadas en términos de las funciones de escala $\phi_{m,k}(x)$ de un solo nivel como base para las funciones admisibles, mientras que las duales $\phi_{m,k}^*(x)$ como las funciones de prueba. El estudio de la convergencia se realiza a través del análisis de Fourier. Para esquemas de Galerkin, en donde $\phi = \phi^*$, este procedimiento se ha aplicado en el método de elementos finitos y aproximación spline.

4.2. Método de Petrov-Galerkin para la ecuación KdV

El método de Petrov-Galerkin está inmerso en un método más general, como lo es el Método de los Residuos Ponderados. Comencemos entonces presentando algunos hechos que muestran de donde aparecen las ideas fundamentales del método.

Sea V un espacio de Hilbert separable y H un subespacio denso de V . Si existe $u \in V$ tal que $(u, v) = 0$ para todo $v \in H$, entonces $u = 0$ en V , donde $(\cdot, \cdot)$ es el producto interno en V . Supongamos ahora que $\{\psi_i\}$ es una base en V , entonces $(u, \psi_k) = 0$ para todo k , esto implica que $u = 0$ en V .

Consideremos el operador $T : D_T \subset V \rightarrow V$, tal que

$$Tu = f \quad \text{en } \Omega \quad (4.2.1)$$

con condiciones de frontera homogéneas apropiadas. Los elementos de D_T satisfacen esas condiciones de frontera. Si $u \in D_T$ es tal que

$$(Tu - f, \psi_k) = 0 \quad \text{para todo } k = 1, 2, \dots \quad (4.2.2)$$

donde $\{\psi_k\}$ es una base en V , entonces tenemos $Tu - f = 0$ en V , esto es, u es solución de (4.2.1) en V . En otras palabras, encontrar la solución de (4.2.1) es equivalente a encontrar la de (4.2.2). Esta equivalencia es la esencia del método de los residuos ponderados. Observe que u no necesariamente se representa con respecto a la base $\{\psi_i\}$. Cualquier base $\{\phi_i\}$ en D_T se puede usar para representar a u , mientras que el residual, R_N , definido por $R_N = Tu_N - f$ es ortogonal al subespacio generado por $\{\psi_k\}$, es decir,

$$(Tu_N - f, \psi_k) = 0 \quad k = 1, 2, \dots, N. \quad (4.2.3)$$

Esta ecuación se conoce con diferentes nombres dependiendo de la elección de la base $\{\psi_k\}$. En el caso general, en el que se elige $\psi_k \neq \phi_k$ se llama método de Petrov-Galerkin, véase por ejemplo, [4, 11, 12, 17, 57]. Cuando elegimos $\psi_k = \phi_k$, tenemos el bien conocido método de Galerkin.

En el método de los residuos ponderados, buscamos una solución aproximada u_N de la ecuación (4.2.1) de la forma

$$u_N = \sum_{i=1}^N c_i \phi_i$$

donde N es un entero positivo arbitrario pero fijo, y c_i son constantes a determinar. Si el operador T es lineal, entonces la ecuación (4.2.3) se puede expresar como

$$\sum_{i=1}^N (T\phi_i, \psi_k) c_i = (f, \psi_k) \quad k = 1, 2, \dots, N.$$

Consideremos ahora la formulación débil de la ecuación KdV

$$\frac{\partial u}{\partial t} + \mu u \frac{\partial u}{\partial x} + \epsilon \frac{\partial^3 u}{\partial x^3} = 0. \quad (4.2.4)$$

Sea $0 \leq \alpha \leq 3$ y $\beta = 3 - \alpha$, para toda función de prueba v , β -regular, se tiene

$$\int_{\mathbb{R}} v \frac{\partial u}{\partial t} dx + \mu \int_{\mathbb{R}} v u \frac{\partial u}{\partial x} dx + \epsilon \int_{\mathbb{R}} v \frac{\partial^3 u}{\partial x^3} dx = 0,$$

la derivada débil para el tercer término es

$$\int_{\mathbb{R}} v \frac{\partial^3 u}{\partial x^3} dx = (-1)^\beta \int_{\mathbb{R}} \left(\frac{\partial^\alpha u}{\partial x^\alpha} \right) \left(\frac{d^\beta v}{dx^\beta} \right) dx.$$

Luego la formulación débil se puede expresar como

$$\left(\frac{\partial u}{\partial t}, v \right) + \mu \left(u \frac{\partial u}{\partial x}, v \right) + \epsilon (-1)^\beta \left(\frac{\partial^\alpha u}{\partial x^\alpha}, \frac{d^\beta v}{dx^\beta} \right) = 0, \quad (4.2.5)$$

donde $(\cdot, \cdot)$ es el producto interno en $L_2(\mathbb{R})$. Recordemos que una función g es r -regular, si $x^n g^{(s)}(x) \in L_2(\mathbb{R})$ para todo s tal que $0 \leq s \leq r$ y para todo $n \in \mathbb{N}$. O también, si existe una constante $M_{s,n} > 0$ tal que

$|g^{(s)}(x)| \leq M_{s,n}(1+|x|)^{-n}$, para cada $x \in \mathbb{R}$, para todo índice s tal que $0 \leq s \leq r$ y para todo $n \in \mathbb{N}$.

Aplicando el método de Petrov-Galerkin, tomando como funciones admisibles

$$\phi_{h,k}(x) = h^{-1/2}\phi(h^{-1}x - k), \quad k \in \mathbb{Z},$$

donde ϕ es una función de valor real, r -regular, $r \geq 1$, y $h > 0$ es el paso de discretización. Los espacios de aproximación $V_h \subset L_2(\mathbb{R})$ son generados por $\{\phi_{h,k}(x), k \in \mathbb{Z}\}$, y la solución exacta de la ecuación KdV (4.2.4), se aproxima por la expresión

$$u_h(x, t) = \sum_k U_k(t)\phi_{h,k}(x). \quad (4.2.6)$$

Similarmente, las funciones de prueba se toman de la forma $\phi_{h,k}^*(x)$, definidas en términos de una función r^* -regular dual ϕ^* , con $r + r^* \geq 3$.

En la formulación débil (4.2.5) escogemos $\alpha \leq r$ tal que $\beta \leq r^*$. Si reemplazamos u por la solución (4.2.6) y v por cada función de prueba $\phi_{h,l}^*(x)$, entonces

$$\begin{aligned} & \left(\frac{\partial u}{\partial t}, v \right) + \mu \left(u \frac{\partial u}{\partial x}, v \right) + \epsilon (-1)^\beta \left(\frac{\partial^\alpha u}{\partial x^\alpha}, \frac{d^\beta v}{dx^\beta} \right) = \\ & \int_{\mathbb{R}} \frac{\partial}{\partial t} \left(\sum_k \phi_{h,k}(x) U_k(t) \right) \phi_{h,l}^*(x) dx + \\ & + \mu \int_{\mathbb{R}} \left(\sum_s U_s(t) \phi_{hs}(x) \right) \left(\frac{\partial}{\partial x} \sum_k U_k(t) \phi_{hk}(x) \right) (\phi_{h,l}^*(x)) dx + \\ & + (-1)^\beta \epsilon \int_{\mathbb{R}} \left(\frac{\partial^\alpha}{\partial x^\alpha} \sum_k U_k(t) \phi_{hk}(x) \right) \left(\frac{d^\beta}{dx^\beta} \phi_{h,l}^*(x) \right) dx \\ & = h^{-1} \sum_k \int_{\mathbb{R}} \phi(h^{-1}x - k) \phi^*(h^{-1}x - l) \frac{dU_k(t)}{dt} dx + \\ & + \mu h^{-3/2} \sum_s \sum_k \int_{\mathbb{R}} \phi(h^{-1}x - s) \phi^*(h^{-1}x - l) \frac{d}{dx} \phi(h^{-1}x - k) U_s(t) U_k(t) dx + \\ & + h^{-1} \epsilon (-1)^\beta \sum_k \int_{\mathbb{R}} \frac{d^\alpha}{dx^\alpha} \phi(h^{-1}x - k) \frac{d^\beta}{dx^\beta} \phi^*(h^{-1}x - l) U_k(t) dx = 0. \end{aligned}$$

Si ponemos $U_k(t) = U_k$ y hacemos el cambio de variable $y = h^{-1}x - k$, la expresión anterior la podemos escribir como

$$\sum_k a(k) \frac{dU_k}{dt} + \mu h^{-3/2} \sum_s \sum_k b(l, k) U_s U_k + h^{-3} \epsilon \sum_k c(k) U_k = 0,$$

donde

$$\begin{aligned} a(k) &= \int_{\mathbb{R}} \phi(y) \phi^*(y - k) dy \\ b(l, k) &= \int_{\mathbb{R}} \frac{d\phi(y)}{dy} \phi(y - l) \phi^*(y - k) dy \\ c(k) &= (-1)^\beta \int_{\mathbb{R}} \frac{d^\alpha \phi(y)}{dy^\alpha} \frac{d^\beta \phi^*(y - k)}{dy^\beta} dy. \end{aligned}$$

Los coeficientes U_k se determinan del siguiente sistema de ecuaciones diferenciales ordinarias

$$\sum_k a(l - k) \frac{d}{dt} U_k + \mu h^{-3/2} \sum_s \sum_k b(s - k, l - k) U_s U_k + h^{-3} \epsilon \sum_k c(l - k) U_k = 0. \quad (4.2.7)$$

La forma matricial de esta última ecuación es

$$\frac{d}{dt} LU + \mu U^T MU + \epsilon NU = 0 \quad (4.2.8)$$

donde

$$\begin{aligned} U &= (U_k), & L(l, k) &= a(l - k) \\ M(l, k, s) &= h^{-3/2} b(l - k, l - s) \\ N(l, k) &= h^{-3} c(l - k). \end{aligned}$$

Las condiciones iniciales $U_k(0)$, $k \in \mathbb{Z}$, son los coeficientes de $u_h(x, 0) = R_h u_0 \in V_h$, donde R_h es algún esquema de aproximación inicial que se precisará más adelante.

Con un paso de tiempo $\Delta t = t_{n+1} - t_n$ y aplicando la regla del trapecoidal se tiene

$$\frac{dU}{dt} = \frac{U^{n+1} - U^n}{\Delta t},$$

donde $U^n = U(n\Delta t)$ para $n \geq 0$, luego la ecuación (4.2.8) queda

$$L \left[\frac{U^{n+1} - U^n}{\Delta t} \right] + \mu U^T MU + \epsilon NU = 0.$$

Haciendo $G(U) = \mu U^T M U + \epsilon N U$ se tiene

$$L(U^{n+1} - U^n) + \frac{G(U^{n+1}) + G(U^n)}{2} \Delta t = 0, \quad (4.2.9)$$

esta ecuación se resuelve por el método de iteración de Newton.

4.3. La ecuación KdV linealizada

Hasta ahora sólo hemos requerido que las funciones ϕ y ϕ^* tengan cierta regularidad. Sin embargo, para buenos resultados de aproximación, son necesarias condiciones adicionales que las discutiremos en esta sección.

Comencemos entonces resolviendo la ecuación KdV linealizada por el método de la transformada de Fourier. Esto es, aplicaremos la transformada de Fourier a la ecuación

$$u_t + \mu u_x + \epsilon u_{xxx} = 0 \quad (4.3.1)$$

con la misma condición inicial $u(x, 0) = u_0(x)$.

Recordemos que la transformada de Fourier para la función $u(x, t)$ es

$$\mathcal{F}[u(x, t)] = \int_{\mathbb{R}} u(x, t) e^{-i\omega x} dx = \hat{u}(\omega, t).$$

Nótese que esta transformada también es una función del tiempo; es decir, es una transformada de Fourier ordinaria para cada t fijo. De igual manera obtenemos las transformadas de Fourier para las derivadas

$$\mathcal{F}[u_t(x, t)] = \int_{\mathbb{R}} u_t(x, t) e^{-i\omega x} dx = \frac{\partial}{\partial t} \int_{\mathbb{R}} u(x, t) e^{-i\omega x} dx = \hat{u}_t(\omega, t).$$

$$\mathcal{F}\left[\frac{\partial^n u}{\partial x^n}\right] = \int_{\mathbb{R}} \frac{\partial^n u(x, t)}{\partial x^n} e^{-i\omega x} dx = (i\omega)^n \hat{u}(\omega, t), \quad n \in \mathbb{N}$$

suponiendo que u es una función de decrecimiento rápido. En consecuencia, al aplicar la transformada de Fourier a la ecuación (4.3.1) obtenemos

$$\mathcal{F}(u_t) + \mu \mathcal{F}(u_x) + \epsilon \mathcal{F}(u_{xxx}) = 0$$

o lo que es lo mismo $\hat{u}_t + i(\mu\omega - \epsilon\omega^3)\hat{u} = 0$, donde $\hat{u} = \hat{u}(\omega, t)$, esta es una ecuación diferencial de primer orden con coeficientes constantes. Su solución general es

$$\hat{u}(\omega, t) = A(\omega) e^{-it(\mu\omega - \epsilon\omega^3)}.$$

Como $u(x, 0) = u_0(x)$ entonces $\hat{u}(\omega, 0) = \hat{u}_0(\omega)$ por lo tanto, $\hat{u}(\omega, 0) = A(\omega) = \hat{u}_0(\omega)$. Así,

$$\hat{u}(\omega, t) = \hat{u}_0(\omega)e^{-it(\mu\omega - \epsilon\omega^3)}.$$

Ahora bien, $u(x, t)$ se obtiene al aplicar la transformada inversa de Fourier, esto es,

$$\mathcal{F}^{-1}(\hat{u}(\omega, t)) = \mathcal{F}^{-1}\left(\hat{u}_0(\omega)e^{-it(\mu\omega - \epsilon\omega^3)}\right)$$

o

$$u(x, t) = \mathcal{F}^{-1}\left(e^{-it(\mu\omega - \epsilon\omega^3)}\hat{u}_0(\omega)\right).$$

Si definimos el operador lineal y acotado $E(t)$ en $L_2(\mathbb{R})$ por

$$E(t)v := \mathcal{F}^{-1}\left(e^{-it(\mu\omega - \epsilon\omega^3)}\hat{v}(\omega)\right)$$

entonces la solución u se puede escribir como

$$u(x, t) = E(t)u_0(x).$$

Por otro lado, la formulación débil de la ecuación KdV linealizada (4.3.1) es

$$\int_{\mathbb{R}} vu_t dx + \mu \int_{\mathbb{R}} vu_x dx + \epsilon \int_{\mathbb{R}} vu_{xxx} dx = 0, \quad (4.3.2)$$

donde $v \in C_0^\infty(\mathbb{R})$ es una función de prueba β -regular, $0 \leq \alpha \leq 3$ y $\beta = 3 - \alpha$. Como

$$\int_{\mathbb{R}} vu_{xxx} dx = (-1)^\beta \int_{\mathbb{R}} \left(\frac{\partial^\alpha u}{dx^\alpha} \frac{d^\beta v}{dx^\beta} \right) dx,$$

entonces la ecuación (4.3.2) se puede expresar como

$$\left(\frac{\partial u}{\partial t}, v \right) + \mu \left(\frac{\partial u}{\partial x}, v \right) + \epsilon (-1)^\beta \left(\frac{\partial^\alpha u}{dx^\alpha}, \frac{d^\beta v}{dx^\beta} \right) = 0,$$

donde $(\cdot, \cdot)$ es el producto interno en $L_2(\mathbb{R})$.

Como en el caso no lineal, al considerar los espacios $V_h \subset L_2(\mathbb{R})$ con la solución aproximada

$$u_h(x, t) = \sum_k U_k(t) \phi_{hk}(x)$$

y al reemplazar u por u_h y a su vez v por $\phi_{hl}^*(x) = h^{-1/2}\phi^*(h^{-1}x - l)$ se tiene

$$\begin{aligned}
 & \int_{\mathbb{R}} v u_t dx + \mu \int_{\mathbb{R}} v u_x dx + \epsilon(-1)^\beta \int_{\mathbb{R}} \left(\frac{\partial^\alpha u}{\partial x^\alpha} \frac{\partial^\beta v}{\partial x^\beta} \right) dx = \\
 & \int_{\mathbb{R}} (h^{-1/2}\phi^*(h^{-1}x - l)) \left(\frac{\partial}{\partial t} \sum_k U_k(t) \phi_{hk}(x) \right) dx + \\
 & + \mu \int_{\mathbb{R}} \left(h^{-1/2}\phi^*(h^{-1}x - l) \right) \left(\frac{\partial}{\partial x} \sum_k U_k(t) \phi_{hk}(x) \right) dx + \\
 & + \epsilon(-1)^\beta \int_{\mathbb{R}} \frac{\partial^\alpha}{\partial x^\alpha} \left(\sum_k U_k(t) \phi_{hk}(x) \right) \frac{\partial^\beta}{\partial x^\beta} (h^{-1/2}\phi^*(h^{-1}x - l)) dx \\
 & = \int_{\mathbb{R}} h^{-1/2}\phi(h^{-1}x - l) \sum_k (h^{-1/2}\phi^*(h^{-1}x - l)) \frac{d}{dt} U_k(t) dx + \\
 & + \mu h^{-1/2} \int_{\mathbb{R}} \phi^*(h^{-1}x - l) \sum_k U_k \frac{d}{dx} (h^{-1/2}\phi(h^{-1}x - k)) dx + \\
 & + \epsilon(-1)^\beta \int_{\mathbb{R}} h^{-1/2} \sum_k U_k(t) \frac{d^\alpha}{dx^\alpha} (h^{-1/2}\phi(h^{-1}x - k)) h^{-1/2} \frac{d^\beta}{dx^\beta} \phi^*(h^{-1}x - l) dx \\
 & = h^{-1} \sum_k \int_{\mathbb{R}} \phi^*(h^{-1}x - l) \phi(h^{-1}x - k) \frac{dU_k(t)}{dt} dx + \\
 & + \mu h^{-1} \sum_k \int_{\mathbb{R}} \phi^*(h^{-1}x - l) \frac{d}{dx} \phi(h^{-1}x - k) U_k(t) dx \\
 & + \epsilon(-1)^\beta h^{-1} \sum_k \int_{\mathbb{R}} \frac{d^\alpha}{dx^\alpha} \phi(h^{-1}x - k) \frac{d^\beta}{dx^\beta} \phi^*(h^{-1}x - l) U_k(t) dx = 0
 \end{aligned}$$

con el cambio de variable $y = h^{-1}x - k$ y poniendo $U_k(t) = U_k$ se tiene

$$\sum_k a(l - k) \frac{dU_k}{dt} + h^{-1} \mu \sum_k d(l - k) U_k + \epsilon(-1)^\beta h^{-3} \sum_k c(l - k) U_k = 0$$

que finalmente se puede escribir

$$\sum_k \left[a(l - k) \frac{dU_k}{dt} + h^{-1} [\mu d(l - k) + h^{-2} \epsilon c(l - k)] U_k \right] = 0, \quad (4.3.3)$$

donde

$$d(k) = \int_{\mathbb{R}} \frac{d\phi(x)}{dx} \phi^*(x - k) dx.$$

De manera análoga se obtiene el sistema equivalente a (4.2.9)

$$\sum_k a(l-k) [U_k^{n+1} - U_k^n] + h^{-1} \Delta t \sum_k [\mu d(l-k) + h^{-2} \epsilon c(l-k)] \frac{U_k^{n+1} + U_k^n}{2} = 0. \quad (4.3.4)$$

Las ecuaciones (4.3.3) y (4.3.4) están en la forma de convolución discreta, luego aplicando la transformada discreta de Fourier

$$\tilde{a}(\omega) = \sum_{k \in \mathbb{Z}} a(k) e^{-ik\omega} = \sum_{k \in \mathbb{Z}} a_k e^{-ik\omega},$$

donde $a = (\dots, a_{-1}, a_0, a_1, \dots) \in \ell^2(\mathbb{Z})$, se tiene respectivamente

$$\tilde{a}(\omega) \frac{d}{dt} \tilde{U}(\omega, t) + h^{-1} [\mu \tilde{d}(\omega) + h^{-2} \epsilon \tilde{c}(\omega)] \tilde{U}(\omega, t) = 0$$

$$\tilde{a}(\omega) [\tilde{U}^{n+1}(\omega) - \tilde{U}^n(\omega)] + h^{-1} \Delta t [\mu \tilde{d}(\omega) + h^{-2} \epsilon \tilde{c}(\omega)] \frac{\tilde{U}^{n+1}(\omega) + \tilde{U}^n(\omega)}{2} = 0,$$

la primera de estas ecuaciones se puede escribir como

$$\frac{d}{dt} \tilde{U}(\omega, t) + h^{-1} \left[\frac{\mu \tilde{d}(\omega) + h^{-2} \epsilon \tilde{c}(\omega)}{\tilde{a}(\omega)} \right] \tilde{U}(\omega, t) = 0$$

o de forma más corta

$$\frac{d}{dt} \tilde{U}(\omega, t) + \frac{Q_h(\omega)}{h} \tilde{U}(\omega, t) = 0$$

donde

$$Q_h(\omega) = \frac{\mu \tilde{d}(\omega) + h^{-2} \epsilon \tilde{c}(\omega)}{\tilde{a}(\omega)}$$

la solución de la ecuación diferencial es

$$\tilde{U}(\omega, t) = c e^{-\left(\frac{Q_h(\omega)t}{h}\right)}$$

la condición inicial para $t = 0$ es $\tilde{U}(\omega, 0) = c$ luego

$$\tilde{U}(\omega, t) = \tilde{U}(\omega, 0) e^{-\left(\frac{Q_h(\omega)t}{h}\right)}.$$

Para la otra ecuación tenemos

$$\tilde{U}^{n+1}(\omega) - \tilde{U}^n(\omega) + h^{-1} \Delta t \left[\frac{\mu \tilde{d}(\omega) + h^{-2} \epsilon \tilde{c}(\omega)}{\tilde{a}(\omega)} \right] \frac{\tilde{U}^{n+1}(\omega) + \tilde{U}^n(\omega)}{2} = 0$$

y agrupando términos se tiene

$$\tilde{U}^{n+1}(\omega) \left[1 + \Delta t \frac{Q_h(\omega)}{2h} \right] - \tilde{U}^n(\omega) \left[1 - \Delta t \frac{Q_h(\omega)}{2h} \right] = 0$$

luego la solución de la ecuación en diferencias con el valor inicial dado es

$$\tilde{U}^n(\omega) = \left[\frac{1 - \left(\frac{\Delta t}{2h}\right) Q_h(\omega)}{1 + \left(\frac{\Delta t}{2h}\right) Q_h(\omega)} \right]^n \tilde{U}(\omega, 0)$$

o de forma más concisa $\tilde{U}^n(\omega) = [\lambda_h(\omega)]^n \tilde{U}(\omega, 0)$, donde

$$\lambda_h(\omega) = \left[\frac{1 - \left(\frac{\Delta t}{2h}\right) Q_h(\omega)}{1 + \left(\frac{\Delta t}{2h}\right) Q_h(\omega)} \right].$$

La solución $u_h(x, t) = \sum_k U_k(t) \phi_{hk}(x)$ tiene transformada de Fourier dada por

$$\hat{u}_h(\omega, t) = \sum_k U_k(t) h^{-1/2} \int_{\mathbb{R}} \phi(h^{-1}x - k) e^{-ix\omega} dx, \quad k \in \mathbb{Z}$$

con el cambio de variable $y = h^{-1}x - k$ tenemos

$$\begin{aligned} \hat{u}_h(\omega, t) &= \sum_k U_k(t) h^{1/2} \int_{\mathbb{R}} e^{-i(y+k)h\omega} \phi(y) dy \\ &= h^{1/2} \sum_k U_k(t) e^{-ikh\omega} \int_{\mathbb{R}} e^{-iyh\omega} \phi(y) dy \\ &= h^{1/2} \tilde{U}(h\omega, t) \hat{\phi}(h\omega) \\ &= \tilde{U}(h\omega, 0) e^{-\left(\frac{Q_h(h\omega)t}{h}\right)} h^{1/2} \hat{\phi}(h\omega) \end{aligned}$$

como $u_h(x, 0) = R_h u_0(x)$, entonces

$$\hat{u}_h(\omega, 0) = \mathcal{F}(R_h u_0)(\omega) = \tilde{U}(h\omega, 0) h^{1/2} \hat{\phi}(h\omega)$$

por lo tanto,

$$\hat{u}_h(\omega, t) = e^{-\left(\frac{Q_h(h\omega)t}{h}\right)} \mathcal{F}(R_h u_0)(\omega).$$

Aplicando transformada inversa de Fourier tenemos

$$u_h(x, t) = \mathcal{F}^{-1} \left[e^{-\left(\frac{Q_h(h\omega)t}{h}\right)} \mathcal{F}(R_h u_0)(\omega) \right].$$

Si definimos el operador $F_h(t)$ por

$$F_h(t)v = \mathcal{F}^{-1}\left[e^{-\left(\frac{Q_h(h\omega)t}{h}\right)}\hat{v}\right]$$

entonces la solución $u_h(x, t)$ se puede expresar como

$$u_h(x, t) = F_h(t)[R_h u_0(x)].$$

Observe que $F_h(t)v$ se puede escribir en la forma de convolución discreta, esto es,

$$F_h(t)v(x) = \sum_k f_k\left(\frac{t}{h}\right)v(x - kh),$$

donde $f_k(t/h)$ son los coeficientes de Fourier de la exponencial $e^{-Q_h(\omega)t/h}$. De igual manera, la solución discreta

$$u_h^n(x) = \sum_k U_k^n \phi_{hk}(x)$$

tiene transformada de Fourier

$$\hat{u}_h^n(\omega) = \sum_k h^{-1/2} U_k^n \int_{\mathbb{R}} \phi(h^{-1}x - k) e^{-ix\omega} dx$$

nuevamente haciendo $y = h^{-1}x - k$ se tiene

$$\begin{aligned} \hat{u}_h^n(\omega) &= \sum_k U_k^n e^{-i\omega kh} h^{1/2} \int_{\mathbb{R}} \phi(y) e^{-ih\omega y} dy \\ &= \tilde{U}^n(h\omega) h^{1/2} \hat{\phi}(h\omega), \end{aligned}$$

pero $\tilde{U}^n = \tilde{U}(\omega, 0)[\lambda_n(\omega)]^n$, así,

$$\begin{aligned} \hat{u}_h^n(\omega) &= \tilde{U}(h\omega, 0)[\lambda_n(h\omega)]^n h^{1/2} \hat{\phi}(h\omega) \\ &= [\lambda_n(h\omega)]^n \mathcal{F}(R_h u_0)(\omega) \end{aligned}$$

al aplicar la transformada inversa de Fourier se tiene

$$\begin{aligned} u_h^n(x) &= \mathcal{F}^{-1}\left([\lambda_n(h\omega)]^n \mathcal{F}(R_h u_0)(\omega)\right) \\ &= G_h^n(R_h u_0(x)), \end{aligned}$$

donde $G_h^n v = \mathcal{F}^{-1}([\lambda_n(h\omega)]^n \hat{v})$.

Nótese que el método es estable si $\operatorname{Re}(Q_h(\omega)) \geq 0$ para cada ω real, acá $\operatorname{Re}(z)$ es la parte real del número complejo z . El método se llama conservativo si $\operatorname{Re}(Q_h(\omega)) = 0$ para todo ω , o disipativo si $\operatorname{Re}(Q_h(\omega)) > 0$ sobre algún intervalo.

4.4. Convergencia y estabilidad

Estudiaremos la convergencia puntual de las soluciones aproximadas $u_h(x, t)$ y $u_h^n(x)$ en los puntos de malla $x = hk$. Con el propósito de evitar errores debido a la aproximación del dato inicial, supondremos que $R_h u_0$ interpola u_0 en los puntos de malla.

También se supondrá que ϕ es r -regular, $\hat{\phi}(0) \neq 0$ y $\hat{\phi}(\omega)$ tiene ceros de orden $p+1$ para todos los puntos $\omega = 2k\pi$, $k \in \mathbb{Z}$ no nulo, para algún entero $p \geq 0$. El conjunto de todas las funciones que satisfacen estas propiedades, se denota por $\mathcal{H}_{r,p}$.

4.4.1. Resultados de convergencia

Comencemos entonces considerando la siguiente condición de interpolación

$$\sum_k \phi(k) e^{-ik\omega} = \sum_k \hat{\phi}(\omega + 2k\pi) \neq 0, \quad (4.4.1)$$

para todo ω real, con el fin de definir un operador de interpolación en términos de las trasladadas y dilatadas de la función ϕ . Como el operador de interpolación R_h conmuta con $F_h(t)$, como también con G_h^n , entonces la solución aproximada $u_h(x, t)$ y la solución en diferencias finitas $F_h(t)u_0(x, t)$ (respectivamente $u_h^n(x)$ y $G_h^n u_0(x)$) coinciden sobre los puntos de la malla.

Supongamos ahora que $\phi \in \mathcal{H}_{r,p}$ y $\phi^* \in \mathcal{H}_{r^*,p^*}$. Para $0 \leq \alpha \leq r$ y $0 \leq \beta \leq r^*$ definamos la función 2π -periódica $\zeta_{\alpha,\beta}(\omega) = \sum_{k \in \mathbb{Z}} I_{\alpha,\beta}(k) e^{-ik\omega}$ donde

$$I_{\alpha,\beta}(k) = \int_{-\infty}^{\infty} \frac{d^\alpha \phi}{dx^\alpha}(x) \frac{d^\beta \phi^*}{dx^\beta}(x-k) dx,$$

entonces $\zeta_{\alpha,\beta}(\omega)$ define una función de clase C^∞ y

$$\zeta_{\alpha,\beta}(\omega) = i^{\alpha-\beta} \omega^{\alpha+\beta} \hat{\phi}(\omega) \overline{\hat{\phi}^*(\omega)} + O(\omega^{p+p^*+2}), \quad \text{cuando } \omega \rightarrow 0.$$

En efecto, para aplicar la relación de Parseval, hacemos

$$f(x) = \frac{d^\alpha \phi}{dx^\alpha}(x) \quad \text{y} \quad g(x) = \frac{d^\beta \phi^*}{dx^\beta}(x - k),$$

entonces

$$\begin{aligned} I_{\alpha,\beta}(k) &= \int_{-\infty}^{\infty} \frac{d^\alpha \phi}{dx^\alpha}(x) \frac{d^\beta \phi^*}{dx^\beta}(x - k) dx = \int_{-\infty}^{\infty} f(x) \overline{g(x - k)} dx \\ &= 2\pi \int_{\mathbb{R}} f(x) \left[\int_{\mathbb{R}} \overline{\hat{g}(\omega)} e^{-i\omega(x-k)} d\omega \right] dx \\ &= 2\pi \int_{\mathbb{R}} \overline{\hat{g}(\omega)} e^{i\omega k} \left[\int_{\mathbb{R}} f(x) e^{-i\omega x} dx \right] d\omega \\ &= 2\pi \int_{\mathbb{R}} \overline{\hat{g}(\omega)} \hat{f}(\omega) e^{-i\omega(-k)} d\omega \\ &= 2\pi \mathcal{F}[\hat{f}\hat{g}^*](-k). \end{aligned}$$

Por lo tanto,

$$\zeta_{\alpha,\beta}(\omega) = \sum_k I_{\alpha,\beta}(k) e^{-ik\omega} = \sum_k 2\pi \mathcal{F}[\hat{f}\hat{g}^*](-k) e^{-ik\omega},$$

si hacemos $h(\omega) = (\hat{f}\hat{g}^*)(\omega)$, entonces aplicado la fórmula (1.2.7) tenemos

$$\zeta_{\alpha,\beta}(\omega) = \sum_k \hat{f}(\omega + 2k\pi) \overline{\hat{g}(\omega + 2k\pi)},$$

Lemarié en [46, Lemme 1, p 159] demuestra que la función $\zeta_{\alpha,\beta}(\omega) \in C^\infty$. Ahora bien, de la propiedad (1.2.5) tenemos

$$\mathcal{F}\left[\frac{d^\alpha \phi}{dx^\alpha}\right](\omega) = i^\alpha \omega^\alpha \hat{\phi}(\omega) \quad \text{y} \quad \mathcal{F}\left[\frac{d^\beta \phi^*}{dx^\beta}\right](\omega) = i^\beta \omega^\beta \hat{\phi}^*(\omega)$$

por lo tanto,

$$\begin{aligned} \zeta_{\alpha,\beta}(\omega) &= \sum_{k \in \mathbb{Z}} \hat{f}(\omega + 2k\pi) \overline{\hat{g}(\omega + 2k\pi)} \\ &= \sum_{k \in \mathbb{Z}} \mathcal{F}\left(\frac{d^\alpha \phi}{dx^\alpha}\right)(\omega + 2k\pi) \overline{\mathcal{F}\left(\frac{d^\beta \phi^*}{dx^\beta}\right)(\omega + 2k\pi)} \\ &= i^{\alpha-\beta} \sum_{k \in \mathbb{Z}} (\omega + 2k\pi)^{\alpha+\beta} \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)}. \end{aligned} \quad (4.4.2)$$

Esta última expresión se puede escribir como

$$\zeta_{\alpha,\beta}(\omega) = i^{\alpha-\beta} \left(\omega^{\alpha+\beta} \hat{\phi}(\omega) \overline{\hat{\phi}^*(\omega)} + \sum_{k \neq 0} (\omega + 2k\pi)^{\alpha+\beta} \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)} \right),$$

si hacemos $\mathcal{R}_{\alpha,\beta}(\omega) = \sum_{k \neq 0} (\omega + 2k\pi)^{\alpha+\beta} \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)}$ entonces

$$\zeta_{\alpha,\beta}(\omega) = i^{\alpha-\beta} \left(\omega^{\alpha+\beta} \hat{\phi}(\omega) \overline{\hat{\phi}^*(\omega)} + \mathcal{R}_{\alpha,\beta}(\omega) \right).$$

Como $\hat{\phi}(\omega)$ y $\hat{\phi}^*(\omega)$ tienen ceros de orden $p+1$ y p^*+1 , respectivamente, entonces $\mathcal{R}_{\alpha,\beta}(\omega)$ tiene ceros de orden $p+p^*+2$, y además, es una función de clase C^∞ . Así, $\mathcal{R}_{\alpha,\beta}(\omega) = O(\omega^{p+p^*+2})$ cuando $\omega \rightarrow 0$. En consecuencia,

$$\zeta_{\alpha,\beta}(\omega) = i^{\alpha-\beta} \omega^{\alpha+\beta} \hat{\phi}(\omega) \overline{\hat{\phi}^*(\omega)} + O(\omega^{p+p^*+2}) \quad (4.4.3)$$

cuando $\omega \rightarrow 0$.

Supongamos ahora que $\phi \in \mathcal{H}_{r,p}$ y $\phi^* \in \mathcal{H}_{r^*,p^*}$ satisfacen la condición de interpolación (4.4.1), donde $r \geq 1$ y $r+r^* \geq 3$, como también la condición de estabilidad, $\text{Re}(Q_h(\omega)) \geq 0$ para cada ω real. Entonces, para el dato inicial suave u_0 y para cada $T > 0$, existe una constante $C > 0$, independiente de h , Δt y u_0 , tal que $0 \leq t \leq T$ y $0 \leq n\Delta t \leq T$, se tiene

$$\begin{aligned} \|u(\cdot, t) - u_h(\cdot, t)\|_{2,h} &= \|u(\cdot, t) - F_h(t)u_0(\cdot, t)\|_{2,h} \\ &\leq Ch^{p+p^*-1} \|u_0\|_{H^{p+p^*+2}} \end{aligned} \quad (4.4.4)$$

$$\begin{aligned} \|u(\cdot, n\Delta t) - u_h^n\|_{2,h} &= \|u(\cdot, n\Delta t) - G_h^n u_0\|_{2,h} \\ &\leq C(h^{p+p^*-1} + \Delta t^2) \|u_0\|_{H^{p+p^*+2}}. \end{aligned} \quad (4.4.5)$$

En efecto, de (4.4.2) se tiene

$$\begin{aligned} \zeta_{0,0}(\omega) &= \tilde{a}(\omega) = \sum_k \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)} \\ &= \sum_k \left[\int_{\mathbb{R}} \phi(x) \phi^*(x - k) dx \right] e^{-ik\omega} = \sum_k a(k) e^{-ik\omega} \\ \tilde{d}(\omega) &= \zeta_{1,0}(\omega) = i \sum_k (\omega + 2k\pi) \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)} \end{aligned}$$

$$\tilde{c}(\omega) = (-1)^\beta \zeta_{\alpha,\beta}(\omega) = -i \sum_k (\omega + 2k\pi)^{\alpha+\beta} \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)}.$$

Por lo tanto,

$$\begin{aligned} Q_h(\omega) &= \frac{\mu \tilde{d}(\omega) + h^{-2} \epsilon \tilde{c}(\omega)}{\tilde{a}(\omega)} = \frac{\mu \sum_{k \in \mathbb{Z}} i(\omega + 2k\pi) \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)}}{\sum_{k \in \mathbb{Z}} \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)}} + \\ &+ \frac{h^{-2} \epsilon \sum_k (-i)(\omega + 2k\pi)^{\alpha+\beta} \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)}}{\sum_k \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)}} \\ &= \frac{i \sum_{k \in \mathbb{Z}} [\mu(\omega + 2k\pi) - h^{-2} \epsilon (\omega + 2k\pi)^3] \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)}}{\sum_{k \in \mathbb{Z}} \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)}}. \end{aligned} \quad (4.4.6)$$

Al descomponer las sumas de esta expresión para $k = 0$ y para $k \neq 0$ y en virtud de (4.4.3) se obtiene

$$Q_h(\omega) = \frac{i [\mu\omega - \epsilon h^{-2} \omega^3] \hat{\phi}(\omega) \overline{\hat{\phi}^*(\omega)} + O(\omega^{p+p^*+2}) + h^{-2} O(\omega^{p+p^*+2})}{\hat{\phi}(\omega) \overline{\hat{\phi}^*(\omega)} + O(\omega^{p+p^*+2})}.$$

Esta expresión se puede escribir también como

$$\hat{\phi}(\omega) \overline{\hat{\phi}^*(\omega)} [Q_h(\omega) - i(\mu\omega + \epsilon h^{-2} \omega^3)] + Q_h(\omega) O(\omega^{p+p^*+2}) = O(\omega^{p+p^*+2}) + h^{-2} O(\omega^{p+p^*+2})$$

por propiedades de los desarrollos asintóticos tales como, $Q_h(\omega) O(\omega^{p+p^*+2}) \rightarrow 0$ cuando $\omega \rightarrow 0$, $\hat{\phi}(\omega) \overline{\hat{\phi}^*(\omega)} = 1 + O(\omega^{p+p^*})$ y $O((h\omega)^{p+p^*+2}) = O(h^{p+p^*+2}) O(\omega^{p+p^*+2})$, véase por ejemplo [9], se tiene

$$Q_h(\omega) = i(\mu\omega + \epsilon h^{-2} \omega^3) + O(\omega^{p+p^*+2}) + h^{-2} O(\omega^{p+p^*+2})$$

$$\begin{aligned} Q_h(h\omega) &= i(\mu h\omega + \epsilon h^{-2} (h\omega)^3) + O((h\omega)^{p+p^*+2}) + h^{-2} O((h\omega)^{p+p^*+2}) \\ &= ih\omega(\mu - \epsilon\omega^2) + O(\omega^{p+p^*+2}) [O(h^{p+p^*+2}) + h^{-2} O(h^{p+p^*+2})]. \end{aligned}$$

De otro lado,

$$\begin{aligned} u(x, t) - u_h(x, t) &= E(t)v(x) - F_h(t)v(x) \\ &= \mathcal{F}^{-1} \left[e^{i(\mu\omega - \epsilon\omega^3)t} \hat{v}(\omega) \right] - \mathcal{F}^{-1} \left[e^{-Q_h(h\omega)t/h} \hat{v}(\omega) \right] \\ &= \mathcal{F}^{-1} \left[(e^{i\omega(\mu - \epsilon\omega^2)t} - e^{-Q_h(h\omega)t/h}) \hat{v}(\omega) \right] \\ &= \frac{1}{2\pi} \int_{-\infty}^{\infty} (e^{i\omega(\mu - \epsilon\omega^2)t} - e^{-Q_h(h\omega)t/h}) \hat{v}(\omega) e^{i\omega x} d\omega \end{aligned}$$

Ahora bien,

$$e^{i\omega(\mu-\epsilon\omega^2)t} - e^{-Q_h(h\omega)t/h} = e^{i\omega(\mu-\epsilon\omega^2)t} - e^{-\frac{t}{h}[ih\omega(\mu-\epsilon\omega^2)+O(\omega^{p+p^*+2})[O(h^{p+p^*+2})+h^{-2}O(h^{p+p^*+2})]]}$$

$$\begin{aligned} |e^{i\omega(\mu-\epsilon\omega^2)t} - e^{-Q_h(h\omega)t/h}| &\leq |e^{i\omega(\mu-\epsilon\omega^2)t}| + \\ &\quad + |e^{-i\omega(\mu-\epsilon\omega^2)t}| e^{-tO(\omega^{p+p^*+2})[O(h^{p+p^*+1})+O(h^{p+p^*+1})]}, \end{aligned}$$

pero

$$|e^{i\omega(\mu-\epsilon\omega^2)t}| = 1 \quad \text{y} \quad |e^{-i\omega(\mu-\epsilon\omega^2)t}| = 1,$$

luego

$$\begin{aligned} |e^{i\omega(\mu-\epsilon\omega^2)t} - e^{-Q_h(h\omega)t/h}| &\leq 1 + e^{-tO(\omega^{p+p^*+2})[O(h^{p+p^*+1})+O(h^{p+p^*+1})]} \\ &\leq Ce^{-tO(\omega^{p+p^*+2})O(h^{p+p^*+1})} \\ &\leq Ch^{p+p^*-1}|\omega|^{p+p^*+2}. \end{aligned}$$

Finalmente tenemos

$$\begin{aligned} \|u(x, t) - F_h(x, t)\|_{2,h} &\leq \frac{1}{2\pi} \int_{-\infty}^{\infty} |e^{i\omega(\mu-\epsilon\omega^2)t} - e^{-Q_h(h\omega)t/h}| |\hat{u}_0(\omega)| |e^{i\omega x}| d\omega \\ &\leq Ch^{p+p^*-1} \int_{-\infty}^{\infty} |\omega|^{p+p^*+2} |\hat{u}_0(\omega)| d\omega \\ &\leq Ch^{p+p^*-1} \int_{-\infty}^{\infty} (1 + |\omega|^2)^{p+p^*+2} |\hat{u}_0(\omega)|^2 d\omega \\ &= Ch^{p+p^*-1} \|u_0\|_{H^{p+p^*+2}}. \end{aligned}$$

De manera análoga se prueba la desigualdad (4.4.5).

4.4.2. Condiciones de estabilidad

Comencemos estudiando algunas hipótesis que garanticen las condiciones de estabilidad. Para ello supongamos que $\hat{\phi}(\omega)\hat{\phi}^*(\omega)$ es real para todo $\omega \in \mathbb{R}$. Entonces la ecuación (4.4.6) implica que $\text{Re}(Q_h(\omega)) = 0$, lo que significa que el método es estable y conservativo. Esto sucede si

- $\phi = \phi^*$ (Método de Galerkin);

- ϕ y ϕ^* son simétricas alrededor del mismo punto.

Supongamos que las funciones de prueba se obtienen de $\eta_\lambda(x) = \phi^*(x - \lambda)$, una versión trasladada de ϕ^* , para algún $\lambda > 0$. Como $\hat{\eta}_\lambda(\omega) = e^{-i\lambda\omega} \hat{\phi}^*(\omega)$, entonces $\hat{\eta}_\lambda(\omega + 2k\pi) = e^{-i\lambda(\omega+2k\pi)} \hat{\phi}^*(\omega + 2k\pi) = e^{-i\lambda\omega} e^{-i2k\pi\lambda} \hat{\phi}^*(\omega + 2k\pi)$, luego $\overline{\hat{\eta}_\lambda(\omega + 2k\pi)} = e^{i\lambda\omega} \overline{\hat{\phi}^*(\omega + 2k\pi)} \xi_\lambda$, donde $\xi_\lambda = e^{i2k\lambda\pi}$. En este caso la ecuación (4.4.6) será

$$Q_{h,\lambda}(\omega) = \frac{i \sum_k [\mu(\omega + 2k\pi) - \epsilon h^{-2}(\omega + 2k\pi)^3] \hat{\phi}(\omega + 2k\pi) \overline{e^{i\lambda\omega} \hat{\phi}^*(\omega + 2k\pi)} \xi_\lambda}{\sum_k \hat{\phi}(\omega + 2k\pi) \overline{e^{i\lambda\omega} \hat{\phi}^*(\omega + 2k\pi)} \xi_\lambda},$$

como $e^{i\lambda\omega}$ no depende de k , la última expresión se puede escribir

$$Q_{h,\lambda}(\omega) = \frac{i \sum_k [\mu(\omega + 2k\pi) - \epsilon h^{-2}(\omega + 2k\pi)^3] \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)} \xi_\lambda}{\sum_k \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)} \xi_\lambda},$$

si hacemos $\nu = 1 - \lambda$ entonces

$$\xi_\nu = e^{i2k\nu\pi} = e^{i2k(1-\lambda)\pi} = e^{2k\pi i} e^{-i2k\lambda\pi} = \overline{\xi_\lambda},$$

luego teniendo en cuenta que $\hat{\phi}(\omega) \overline{\hat{\phi}^*(\omega)}$ es real se tiene

$$\overline{Q_{h,\lambda}(\omega)} = - \frac{i \sum_k [\mu(\omega + 2k\pi) - \epsilon h^{-2}(\omega + 2k\pi)^3] \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)} \overline{\xi_\lambda}}{\sum_k \hat{\phi}(\omega + 2k\pi) \overline{\hat{\phi}^*(\omega + 2k\pi)} \overline{\xi_\lambda}},$$

así, $\overline{Q_{h,\lambda}(\omega)} = -Q_{h,\nu}(\omega)$ y esto implica

$$\operatorname{Re}(Q_{h,\nu}(\omega)) = -\operatorname{Re}(Q_{h,\lambda}(\omega)).$$

Esta propiedad se utiliza en las siguientes conclusiones para ϕ y ϕ^* cuando se cumple que $\hat{\phi}(\omega) \overline{\hat{\phi}^*(\omega)}$ es un real

- Para $\lambda = 1/2$, $\operatorname{Re}(Q_{h,\nu}(\omega)) = 0$ para todo real ω , y el método es estable y conservativo.
- Si para algún $0 < \lambda < 1$ el método es estable y conservativo entonces también es verdadero para el parámetro de traslación $1 - \lambda$.
- Si para algún $0 < \lambda < 1$ el método es estable y disipativo, entonces para el parámetro de traslación $1 - \lambda$ es inestable.

4.5. Algoritmos para calcular $a(k)$, $b(l, k)$ y $c(k)$

Para calcular los coeficientes $a(k)$, $b(l, k)$ y $c(k)$ recurriremos a resolver (parcialmente) el problema de valor de frontera periódico para la ecuación diferencial con coeficientes constantes

$$\sum_{s=0}^d a_s D^s u \equiv \mathcal{P}u = g$$

u y g 1 – periódicas en el espacio

donde $D = \frac{d}{dx}$. El método de Galerkin descrito en esta sección se basa en el marco biortogonal $\{V_j, V_j^*\}$, asociado a las funciones de escala spline biortogonal $\phi = \phi_N$ y $\phi^* = \phi_{N,N^*}$, con el correspondiente par conjugado $\{\tilde{\phi}, \tilde{\phi}^*\}$, donde

$$\tilde{\phi}(x) = \begin{cases} \phi_{N-1}(x) & \text{si } N \text{ es impar,} \\ \phi_{N-1}(x+1) & \text{si } N \text{ es par.} \end{cases}$$

y

$$\tilde{\phi}^*(x) = \begin{cases} \phi_{N-1,N^*+1}(x) & \text{si } N \text{ es impar,} \\ \phi_{N-1,N^*+1}(x+1) & \text{si } N \text{ es par.} \end{cases}$$

Una de las familias de funciones base se utilizará como funciones de ensayo o admisibles, digamos $\phi_{j,k}(x)$, y su familia dual $\phi_{j,k}^*(x)$ como funciones de prueba. Esto significa que la solución aproximada para el problema de valor de frontera dado es

$$u_j(x) = \sum_{k \in \mathbb{Z}} U_j(k) \phi_{j,k}(x),$$

donde $U_j(k) = U_j(k + 2^j)$ son los coeficientes periódicos que se deben determinar a través de la siguiente relación

$$\int_{\mathbb{R}} (\mathcal{P}u_j - g) \phi_{j,l}^*(y) dy = 0 \quad \text{para cada } l \in \mathbb{Z}. \quad (4.5.1)$$

Siempre que las funciones base tengan la suficiente regularidad para que esta integral tenga sentido. La ecuación (4.5.1) se puede expresar como un sistema de ecuaciones diferenciales ordinarias

$$\sum_{k \in \mathbb{Z}} \sum_{s=0}^d 2^{sj} \Gamma(k-l) U_j(k) = G_j(l) \quad \text{para cada } l \in \mathbb{Z}, \quad (4.5.2)$$

donde $G_j(k) = \langle g, \psi_{j,k}^* \rangle$ y

$$\Gamma_s(m) = \Gamma_{s,N,N^*}(m) = (-1)^p \int_{\mathbb{R}} D^{s-p} \phi(y) D^p \phi^*(y+m) dy, \quad (4.5.3)$$

los índices $0 \leq p \leq s$ se introducen aquí con el fin de controlar la posible falta de regularidad de las funciones base. Para ilustrar el cálculo de los coeficientes utilizando la fórmula (4.5.3), lo haremos para la función hat, con $0 \leq s \leq 3$,

$$\phi(x) = \phi_2(x) = \begin{cases} 1+x & \text{si } -1 \leq x < 0 \\ 1-x & \text{si } 0 \leq x \leq 1 \\ 0 & \text{en otro caso} \end{cases}$$

y sus duales $\phi^* = \phi_{2,N^*}$ con $N^* \geq 4$, que tienen soporte sobre los intervalos $[-N^*, N^*]$. De la relación de biortogonalidad se tiene

$$\Gamma_0(m) = a(k) = \int_{\mathbb{R}} \phi(y) \phi^*(y) dy = \delta_{k,0}.$$

Si derivamos la función hat obtenemos

$$\frac{d\phi_2(x)}{dx} = \begin{cases} 1 & -1 < x < 0, \\ -1 & 0 < x < 1 \\ 0 & \text{en otro caso} \end{cases}$$

entonces

$$d(-m) = \Gamma_1(m) = \int_{\mathbb{R}} \frac{d\phi(y)}{dy} \phi^*(y+m) dy = \int_{-1}^0 \phi^*(y+m) dy - \int_0^1 \phi^*(y+m) dy$$

para $s = 1$, $p = 0$ y $N = 2$

$$\tilde{\phi}^*(x) = \tilde{\phi}_{2,N^*}(x) = \phi_{1,N^*+1}(x+1) = \int_x^{x+1} \phi_{2,N^*}(y) dy$$

aplicando estos cálculos obtenemos

$$\begin{aligned} \Gamma_1(m) &= \int_{-1}^0 \phi_{2,N^*}(y+m) dy - \int_0^1 \phi_{2,N^*}(y+m) dy \\ &= \int_{-1+m}^m \phi_{2,N^*}(z) dz - \int_m^{1+m} \phi_{2,N^*}(z) dz \\ &= \phi_{1,N^*+1}(m) - \phi_{1,N^*+1}(m+1). \end{aligned}$$

La segunda derivada de la función hat en el sentido de las distribuciones es

$$\begin{aligned}
 \langle \phi_2''(x), \varphi \rangle &= (-1)^2 \langle \phi_2(x), \varphi''(x) \rangle = \int_{-\infty}^{\infty} \phi_2(x) \varphi''(x) dx \\
 &= (x+1) \varphi'(x)|_{-1}^0 + (1-x) \varphi'(x)|_0^1 \\
 &= \varphi'(0) - \varphi(x)|_{-1}^0 - \varphi'(0) + \varphi(x)|_0^1 \\
 &= -\varphi(0) + \varphi(-1) + \varphi(1) - \varphi(0) \\
 &= \delta(x+1) - 2\delta(x) + \delta(x-1)
 \end{aligned}$$

con $s = 2$ y $p = 1$ tenemos

$$\begin{aligned}
 \Gamma_2 &= (-1) \int_{\mathbb{R}} \phi'(y) \phi^{*'}(y+m) dy = - \int_{-1}^0 \phi^{*'}(y+m) dy + \int_0^1 \phi^{*'}(y+m) dy \\
 &= -\phi^*(m) + \phi^*(-1+m) + \phi^*(1+m) - \phi^*(m) \\
 &= \phi_{2,N^*}(m-1) - 2\phi_{2,N^*}(m) + \phi_{2,N^*}(1+m)
 \end{aligned}$$

Nótese que $\phi_{2,N^*}(x) = \int_x^{x+1} \phi_{3,N^*-1}(y) dy$, y con $N = 3$

$$\tilde{\phi}^*(x) = \phi_{2,N^*+1}(x) = \int_x^{x+1} \phi^*(y) dy = \int_x^{x+1} \phi_{3,N^*}(y) dy.$$

Para el cálculo de la integral $c(k)$ tomamos los valores $s = 3$ y $p = 1$, también es válido para valores de $\alpha = 2$ y $\beta = 1$

$$\begin{aligned}
 \Gamma_3(m) &= (-1) \int_{\mathbb{R}} \phi''(y) \phi^{*'}(y+m) dy \\
 &= - \int_{\mathbb{R}} [\delta(y+1) - 2\delta(y) + \delta(y-1)] [\phi^{*'}(y+m)] dy \\
 &= \int_{\mathbb{R}} [\delta(y+1) - 2\delta(y) + \delta(y-1)] [\phi_{3,N^*-1}(y+m+1) - \phi_{3,N^*-1}(y+m)] dy \\
 &= - \int_{\mathbb{R}} [\delta(y+1) \phi_{3,N^*-1}(y+m+1) - 2\delta(y) \phi_{3,N^*-1}(y+m+1) + \\
 &\quad + \delta(y-1) \phi_{3,N^*-1}(y+m+1) - \delta(y+1) \phi_{3,N^*-1}(y+m) + \\
 &\quad + 2\delta(y) \phi_{3,N^*-1}(y+m) - \delta(y-1) \phi_{3,N^*-1}(y+m)] dy \\
 &= -[\phi_{3,N^*-1}(m) - 2\phi_{3,N^*-1}(m+1) + \phi_{3,N^*-1}(m+2) - \phi_{3,N^*-1}(m-1) + \\
 &\quad + 2\phi_{3,N^*-1}(m) - \phi_{3,N^*-1}(m+1)] \\
 &= \phi_{3,N^*-1}(m-1) + 3\phi_{3,N^*-1}(m+1) - 3\phi_{3,N^*-1}(m) - \phi_{3,N^*-1}(m+2)
 \end{aligned}$$

luego

$$c(k) = \Gamma_3(k) = \phi_{3,N^*-1}(k-1) + 3\phi_{3,N^*-1}(k+1) - 3\phi_{3,N^*-1}(k) - \phi_{3,N^*-1}(k+2).$$

Si bien, para el cálculo de la integral $b(l, k)$ no seguimos el mismo procedimiento, si utilizamos algunos resultados obtenidos anteriormente

$$\begin{aligned} b(l, k) &= \int_{-\infty}^{+\infty} \frac{d\phi_2(x)}{dx} \phi_2(x-l) \phi^*(x-k) dx \\ &= \int_{-1}^0 \phi_2(x-l) \phi_{2,N^*}(x-k) dx - \int_0^1 \phi_2(x-l) \phi_{2,N^*}(x-k) dx \\ &= \int_{-1-k}^{-k} \phi_2(y+k-l) \phi_{2,N^*}(y) dy - \int_{-k}^{1-k} \phi_2(y+k-l) \phi_{2,N^*}(y) dy \\ &= \int_{l-1-k}^{l-k} (1+y+k-l) \phi_{2,N^*}(y) dy - \int_{l-k}^{1+l-k} (1-y-k+l) \phi_{2,N^*}(y) dy \\ &= (1+k-l) \int_{-1-k+l}^{-k+l} \phi_{2,N^*}(y) dy + \int_{-1-k+l}^{-k+l} y \phi_{2,N^*}(y) dy + \\ &\quad + \int_{-k+l}^{-1-k+l} y \phi_{2,N^*}(y) dy - (1-k+l) \int_{-k+l}^{-1-k+l} \phi_{2,N^*}(y) dy \end{aligned}$$

por lo tanto,

$$b(l, k) = I(-1-k+l) + I(-k+l) + (1+k-l) \phi_{1,N^*+1}(l-k) - (1-k+l) \phi_{1,N^*+1}(1-k+l),$$

donde

$$I(k) = \int_k^{k+1} y \phi_{2,N^*}(y) dy.$$

Para $l = 0$;

$$b(0, k) = I(-1-k) + I(-k) + (1+k) \phi_{1,N^*+1}(-k) - (1-k) \phi_{1,N^*+1}(1-k).$$

Para $l = -1$

$$b(-1, k) = I(-2-k) + I(-1-k) + (2+k) \phi_{1,N^*+1}(-1-k) + k \phi_{1,N^*+1}(-k).$$

Para $l = 1$,

$$b(1, k) = I(-k) + I(1-k) + k \phi_{1,N^*+1}(1-k) - (2-k) \phi_{1,N^*+1}(2-k).$$

Nótese que $I(1-k) + (k-2)\phi_{1,N^*+1}(2-k) = 0$. Luego

$$b(1, k) = I(-k) + k\phi_{1,N^*+1}(1-k).$$

De igual forma se tiene para $l = -1$ que

$$b(-1, k) = I(-1-k) + k\phi_{1,N^*+1}(-k).$$

Por simetría de $\phi_{2,N^*}(y)$ alrededor de $x = 0$ se tiene

$$\begin{aligned} I(k) &= \int_k^{k+1} \phi_{2,N^*}(y)dy = - \int_{k+1}^k y\phi_{2,N^*}(y)dy = - \int_{-k-1}^{-k} y\phi_{2,N^*}(y)dy \\ &= -I(-k-1) \end{aligned}$$

y la simetría de $\phi_{1,N^*-1}(x)$ alrededor de $x = \frac{1}{2}$, implica

$$b(-1, k) = -b(1, -k) \quad \text{y} \quad b(0, k) = -b(0, -k).$$

De la condición de biortogonalidad aplicada a $\phi = \phi_2$ y $\phi^* = \phi_{2,N^*}$ se obtiene la siguiente relación

$$\int_{-\infty}^{+\infty} \phi(y)\phi^*(y-k)dy = \delta_{k,0}$$

de esta manera se obtiene

$$\begin{aligned} \int_{-\infty}^{+\infty} \phi_2(y)\phi_{2,N^*}(y-k)dy &= \int_{-1-k}^{-k} (1+z+k)\phi_{2,N^*}(z)dz + \\ &+ \int_{-k}^{1-k} (1-z-k)\phi_{2,N^*}(z)dz \\ &= \int_{-1-k}^{-k} z\phi_{2,N^*}(z)dz - \int_{-k}^{1-k} z\phi_{2,N^*}(z)dz + \\ &+ (1+k) \int_{-1-k}^{-k} \phi_{2,N^*}(z)dz + \\ &+ (1-k) \int_{-k}^{1-k} \phi_{2,N^*}(z)dz = \delta_{k,0}. \end{aligned}$$

Luego

$$I(-1-k) - I(-k) + (k+1)\phi_{1,N^*+1}(-k) + (1-k)\phi_{1,N^*+1}(1-k) = \delta_{k,0}$$

esta fórmula se puede usar para simplificar la expresión $b(0, k)$ dando

$$b(0, k) = \delta_{0,k} + 2I(-k) - 2(1 - k)\phi_{1,N^*+1}(k).$$

El algoritmo para el cálculo de $I(k)$ es el siguiente:

- $I(0) = \phi_{1,N^*+1}(0) - \frac{1}{2}$
- Para $k = 1, 2, \dots, N^* - 1$
 - $I(-k) = -I(k - 1)$
 - $I(k) = -I(-k) + (k + 1)\phi_{1,N^*+1}(-k) + (1 - k)\phi_{1,N^*+1}(k)$
- $I(-N^*) = -I(N^* - 1).$

Para llevar a cabo los cálculos es necesario tener en cuenta los valores de ϕ_{1,N^*+1} y ϕ_{3,N^*+1} en los enteros. En general, el orden para calcular los valores de la función escala de soporte compacto en los enteros, puede usarse el algoritmo de cascada, véase la Sección 2.6, que empieza con la ecuación de dilatación y soluciona el problema de valores propios para la matriz cuyas entradas son $h(2l - k)$.

4.6. Implementación del método y resultados numéricos

```
% Implementación del Método Petrov-Galerkin basado
% en Wavelets Biortogonales, para una función hat.

% Se limpian las variables que hay en el workspace
clc; clear all;

% Parámetros de la función KdV
mu=1; epsilon=0.000484; C=0.3; K=(C/(4*epsilon))^0.5; d=-K; xa=0;
xb=2;

% Parámetros del método (ancho de malla, delta de tiempo, etc)
h=1/2^5; delta_t=0.1; x=0:h:2; n_ast=4; tam=length(x);
wname='bior2.4'; % Se trabajo con n_ast=4

% Evaluación de las Wavelets Biortogonales a trabajar, se usó el algoritmo
% biphivals para calcular la wavelet, dicho código se puede encontrar en
% la página de Matlab.
[s9,a9]= biorwavf('bior3.3'); % Wavelet Biortogonal(3,n_ast-1)
[h0,h1,f0,f1] = biorfilt(a9, s9); [x1,phi9,wp1,psi9,psitilde9] =
biphivals(h0,h1,f0,f1,0); x1=x1-3;
```

88 Método Wavelet-Petrov-Galerkin para la ecuación KdV

```
[s9,a9]= biorwavf('bior1.5'); % Wavelet Biortogonal(1,n_ast+1)
[h0,h1,f0,f1] = biorfilt(a9, s9); [x2,phi9,wp2,psi9,psitilde9] =
biphivals(h0,h1,f0,f1,0); x2=x2-4;

% Función a trabajar
phi=@(x) hat Function(x);

% Vector que posee el cálculo de la integral I para el cálculo de B
I = cal_I(n_ast,x2,wp2);

% Se renombran los nombre de las funciones para la evaluación
% de las integrales para el cálculo de coeficientes
a = @(x) x==0; b = @(l,k) calB( l,k, x2, wp2,I, n_ast); c = @(k)
calC(k,x1,wp1);

% Renombro la función que se usa para el método iterativo de Newton.
G=@(U) cal_G(U,h,epsilon,mu,b,c,n_ast );

% Matriz que guarda los coeficientes de la discretización de la malla
% a través del tiempo, una columna i representa los coeficientes en el
% tiempo delta_t*i.
coef_U=zeros(tam,10);

% Hallamos por Mínimos Cuadrados los coeficientes cuando t=0.
y0=3*C*cosh(K*x+d).^2; y0=y0';

R0=zeros(tam); for i=1:tam
    for k=0:tam-1
        R0(i,k+1)=h^(-0.5)*phi(h^(-1)*x(i)-k);
    end
end coef_U(:,1)=(R0'*R0)\R0'*y0;

% Se resuelve el problema a partir de iteración de Newton.
aux=coef_U(:,1); for j=2:10
    i=1;
    while sum((aux-coef_U(:,j)).^2).^5>=1e-7
        D = cal_D( coef_U(:,j),b,c,mu,epsilon,h,delta_t );
        aux=coef_U(:,j);
        coef_U(:,j)=coef_U(:,j)+bicgstabl(D,-(coef_U(:,j)-coef_U(:,j-1)+
        delta_t*(G(coef_U(:,j))+G(coef_U(:,j-1)))/2));
        i=i+1;
    end
    disp('paso');
end

% Calculamos el valor teórico para contrastar con el experimenta
t=0:delta_t:0.9; for j=1:length(t)
    for i=1:length(x)
        Teo(i,j)=3*C*cosh(K*x(i)-K*C*t(j)+d).^2;
    end
end

% Matriz con lo valores de la función calculados numéricamente en los
% diferentes tiempos
Pru=R0*coef_U;
```

4.6 Implementación del método y resultados numéricos 89

```
% Error entre el valor teórico y práctico.
error=sum((Teo-Pru).^2).^5;

% tabla con la información del error en los diferentes tiempos.
tab=[t' error'];

-----

function [b] = calB( l,k, x,wp,I,n_ast)
%CALB Función que calcula el coeficiente b
% l: seria el primer argumento de la función
% k: seria el segundo argumento de la función
% x: vector con las posiciones en las que se cálculo los valores
% de la wavelet
% wp: vector con los valores de la wavelet en los puntos x
% I: un vector con los valores de la integral I
% n_ast: valor con el que se trabaja las wavelet biortogonales

if (k>n_ast || k<=-n_ast)
    b=0;
else
    if (l<=-2 || l>=2)
        b=0;
    elseif (l==-1)
        b=-I(n_ast+1-k-1)-k*encontrarK(-k,x,wp);
    elseif (l==0)
        b=I(n_ast+1-k-1)+I(n_ast+1-k)+(k+1)*encontrar K(-k,x,wp)+(k-1)*encontrar K(-k+1,x,wp);
    elseif (l==1)
        b=-I(n_ast+1-k)-k*encontrarK(-k+1,x,wp);
    end
end

-----

function [c] = calC(k,x1,wp1)
%CALC Función que calcula el coeficiente c
% k: seria el primer argumento de la función
% de la wavelet
% x1: vector con las posiciones en las que se cálculo los valores
% de la wavelet
% wp1: vector con los valores de la wavelet en los puntos x
c=encontrarK(-1-k,x1,wp1)-3*encontrar K(-k,x1,wp1)+ 3*encontrar
K(1-k,x1,wp1)-encontrar K(2-k,x1,wp1); end

-----

function I = cal_I(n_ast, x,wp )
%CAL_I Función que calcula la integral I
% n_ast: valor con el que se trabaja las wavelet biortogonales
% x: vector con las posiciones en las que se cálculo los valores
```

90 Método Wavelet-Petrov-Galerkin para la ecuación KdV

```
%      de la wavelet
%  wp: vector con los valores de la wavelet en los puntos x

I=zeros(n_ast*2+1,1);

I(n_ast+1)=encontrarK(0,x,wp)-1/2; for k=1:n_ast-1
    I(n_ast+1-k)=-I(n_ast+1+k-1);
    I(n_ast+1+k)=-I(n_ast+1-k)+(k+1)*encontrar K(-k,x,wp)+(1-k)*encontrar K(k,x,wp);
end I(1)=-I(2*n_ast);
```

```
-----

function [ G ] = cal_G( U,h,eps,mu,b,c,n_ast )
%CAL_G Función que calcula el operador G
%  U: seria el vector de los coeficientes
%  h: seria el tamaño de la malla
%  eps: parámetro epsilon de la ecuación KdV
%  mu: parámetro mu de la ecuación KdV
%  b: función para el cálculo de b
%  c: función para el cálculo de c
%  n_ast: valor con el que se trabaja las wavelet biortogonales
```

```
G=zeros(length(U),1); for l=0:length(U)-1
    sum1=0;
    sum2=0;
    for k=0:length(U)-1
        if abs(l-k)<=n_ast
            for s=0:length(U)-1
                sum1=sum1+b(s-k,l-k)*U(k+1)*U(s+1);
            end
        end
        sum2=sum2+c(l-k)*U(k+1);
    end
    G(l+1)=mu*h^(-3/2)*sum1+eps*h^(-3)*sum2;
end
```

```
-----

function D = cal_D( U,b,c,mu,epsilon,h,delta_t )
%CAL_D Función que calcula la matriz Jacobiana
%  U: seria el vector de los coeficientes
%  eps: parámetro epsilon de la ecuación KdV
%  mu: parámetro mu de la ecuación KdV
%  b: función para el cálculo de b
%  c: función para el cálculo de c
%  h: seria el tamaño de la malla
%  delta_t: paso en el tiempo.
```

```
D=zeros(length(U));

for l=0:length(U)-1
    for j=0:length(U)-1
        sumas=0;
```

```

    for s=0:length(U)-1
        sumas=sumas+b(s-j,l-j)*U(s+1);
        if s==j
            sumas=sumas+b(s-j,l-j)*U(s+1);
        end
    end
    D(l+1,j+1)=delta_t*(mu*h^(-3/2)*sumas+epsilon*h^(-3)*c(l-j))/2;
end
end

D=D+eye(length(U));

```

```

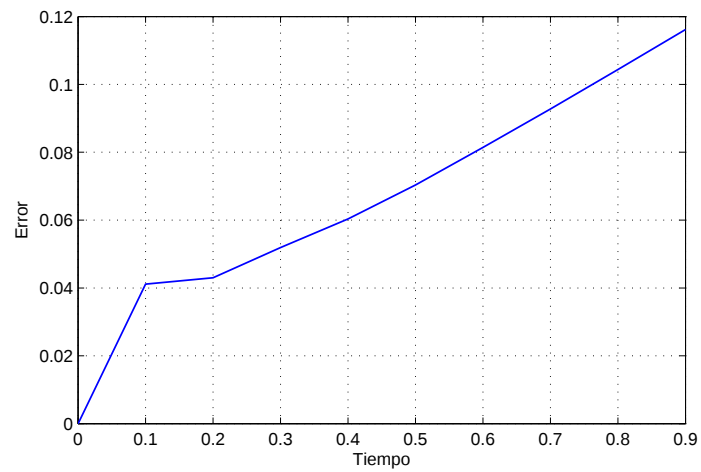
function y0 = encontrar K( x0,x,y )
%ENCONTRAR K Función que encuentra el valor x0 correspondiente a un
%      vector y, buscándolo en un vector x, en caso de no encontrarlo
%      hace una interpolación lineal.
%  x0: valor a encontrar
%  x: vector de búsqueda
%  y: vector con los valores evaluados de x

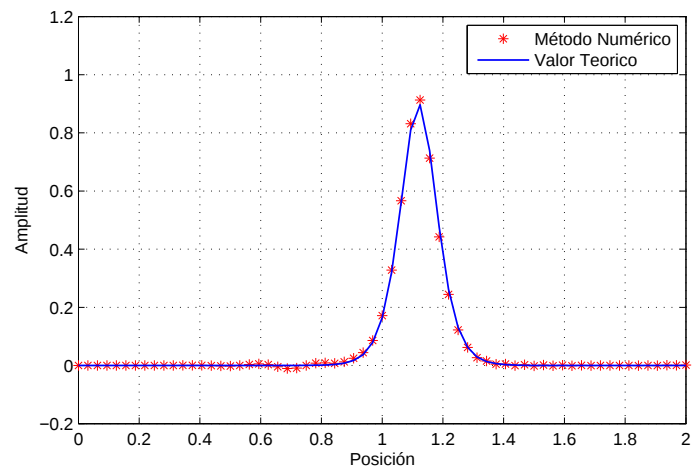
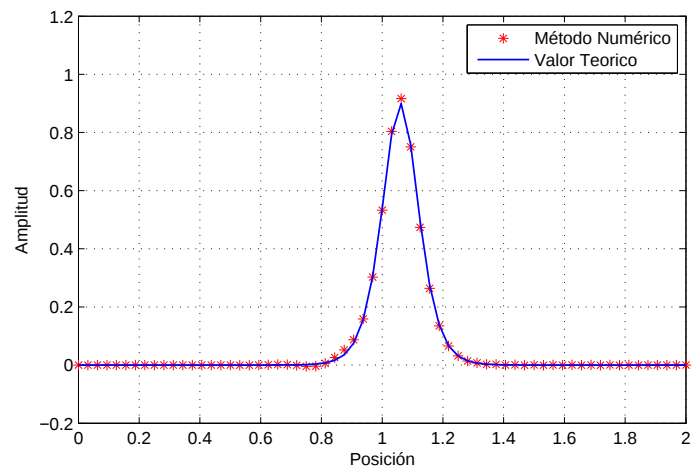
if (x0<=x(1))
    y0=y(1);
elseif (x0>= x(length(x)))
    y0=y(length(y));
else
    for i=1:length(x)-1
        if (x0>=x(i) && x0<=x(i+1))
            lambda=(x0-x(i))/(x(i+1)-x(i));
            y0=lambda*y(i+1)+(1-lambda)*y(i);
        end
    end
end
end

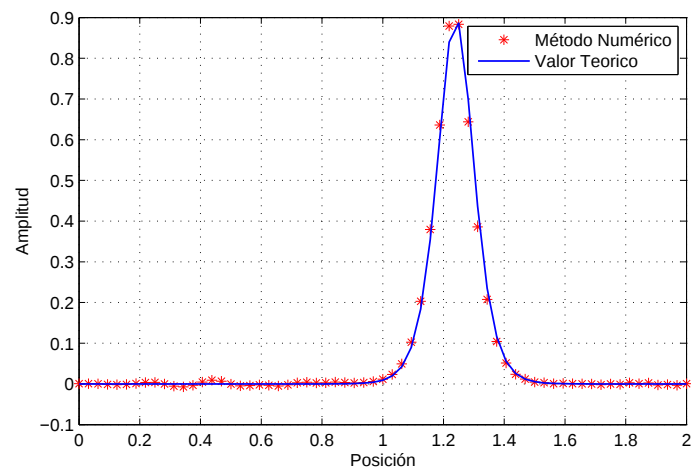
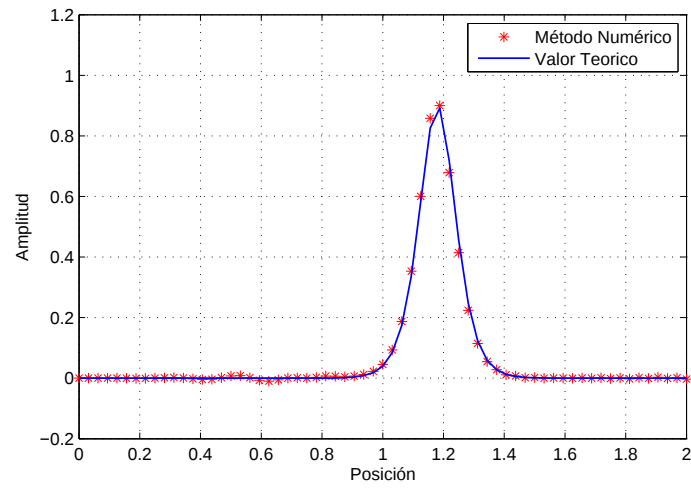
```

Tiempo	Error
0	2,53E-18
0,1	0,04114137
0,2	0,04298284
0,3	0,05190608
0,4	0,06032352
0,5	0,07034124
0,6	0,08141202
0,7	0,09275702
0,8	0,10438704
0,9	0,11619728

Cuadro 4.1: Error del método con $h = 2^{-5}$







Bibliografía

- [1] Adomian, G. *Solving Frontier Problems of Physics: The Decomposition Method*. Kluwer Academic Publishers, Dordrecht, 1994.
- [2] Adomian, G. A new approach to nonlinear partial differential equations, *J. Math. Anal. Appl.*, 102 (1984), 420-434.
- [3] Aldroubi, A. The wavelet transform: A surfing guide, pp 3-36 in *Wavelets in Medicine and Biology*, A. Audroubi, M. Unser (eds.), CRC Press, Boca Raton, Florida, 1996.
- [4] Atkinson, K., and Han, W. *Theoretical Numerical Analysis: A Functional Analysis Framework*. Second Edition. Tex. Appl. Math. 39 Springer-Verlag Berlin-Heidelberg-New York, 2005.
- [5] Beylkin, G., Coifman, R. R., and Rokhlin, V. Fast wavelet transforms and numerical algorithms I, *Commun. Pure and Appl. Math.*, 44 (1991), 141-183.
- [6] Beylkin, G. On wavelet-based algorithms for solving differential equations, pp 449-466 in *Wavelets: Mathematics and Applications*, J. Benedetto, M. Frazier (eds.), CRC Press, Boca Raton, Florida 1994.
- [7] Beylkin, G., Keiser, J. M. On the Adaptive Numerical Solution of Nonlinear Partial Differential Equations in Wavelet bases, *Journal of Computational Physics.*, 132 (1997), 233-259.

-
- [8] Bleistein, N. *Mathematical Methods for Wave Phenomena*. Orlando: Academic Press, 1984.
 - [9] Bleistein, N., Handelsman, R. *Asymptotic Expansions of Integrals*. Dover Publications, Inc. New York, 1986.
 - [10] Boggess, A., Narcowich, F. J. *A First Course in Wavelets with Fourier Analysis*. Prentice Hall, New Jersey, 2001.
 - [11] Braess, D. *Finite Elements: Theory, Fast Solvers, and Applications in Solid Mechanics* 3^a edition. Cambridge University Press. New York, Inc 2007.
 - [12] Brenner, S.C., Scott, L.R. *The Mathematical Theory of Finite Element Methods* 3^a edition. Springer-Verlag: Texts in Applied Mathematics 15. New York, Inc 2008.
 - [13] Burrus, C. S., Gopinath, R. A., Guo, H. *Introduction to Wavelets and Wavelets Transforms A Primer*. Prentice Hall, New Jersey, 1998.
 - [14] Chalub, F. A. C. C., Zubelli, J. P. Sólitos: Na Crista da Onda por mais de 100 Anos, *Matemática Universitária*. No. 30, (2001), 41-65.
 - [15] Chiavassa, G., Liandrat, J. A fully adaptive wavelet algorithm for parabolic partial differential equations, *Applied Numerical Mathematics*. 36 (2001), 333-358.
 - [16] Chui, C. K. *Wavelets: A Mathematical Tool for Signal Analysis*. SIAM Monographs on Mathematical Modeling and Computation, Philadelphia, 1997.
 - [17] Ciarlet, P. G. *The Finite Element Method for Elliptic Problems*. 2nd edition. SIAM: Society for Industrial and Applied Mathematics, Philadelphia, 2002.
 - [18] Cohen, A. *Numerical Analysis of Wavelet Methods*. Volume 32. North-Holland Elsevier Science, Amsterdam, 2003.
 - [19] Cohen, A., Daubechies, I., and Feauveau, J. C. Biorthogonal basis of compactly supported wavelets, *Comm. Pure Appl. Math.* 45 (1992), 485-560.

- [20] Dahlke, S., Weinreich, I. Wavelet-Galerkin methods: an adapted biorthogonal wavelet basis, *Constructive Approximations*. 9(2) (1993), 237-262.
- [21] Dahlke, S., Dahmen, W., Hochmuth, R., Schneider, R. Stable multiscale bases and local error estimation for elliptic problems, *Applied Numerical Mathematics*. 23 (1997), 21-47.
- [22] Dahmen, W., Kunoth, A., and Urban, K. A Wavelet-Galerkin method for the Stokes Equations, *Computing*. 56 (1996), 259-302.
- [23] Dahmen, W. Wavelet methods for PDEs-some recent developments, *Journal of Computational and Applied Mathematics*. 128 (2001), 133-185.
- [24] Daubechies, I. Orthonormal bases of compactly supported wavelets, *Comm. Pure Appl. Math.*, 41 (1988), 909-996.
- [25] Daubechies, I. *Ten Lectures on Wavelets*, CBMS Series 61, SIAM, Philadelphia 1992.
- [26] Daubechies, I. The wavelet transform, time-frequency localization and signal analysis, *IEEE Trans. Inform. Theory*, vol. 36 (1990), 961-1005.
- [27] Debnath, L. *Nonlinear Partial Differential Equations: for Scientists and Engineers*. Second Edition, Birkhäuser, Boston 2005.
- [28] Drazin, P. G., Johnson, R. S. *Solitons: An Introduction*. Second Edition, Cambridge University Press, New York, 1989.
- [29] Ern, A., Guermond, J. L *Theory and Practice of Finite Elements*. Appl. Math. Sci. 159 Springer-Verlag Berlin-Heidelberg-New York, 2010.
- [30] Evans, L. C. *Partial Differential Equations*, AMS, vol. 19, Providence, RI, 1998.
- [31] Fleet, P. V. *Discrete Wavelet Transformations: An Elementary Approach with Applications*, John Wiley & Sons, Inc. New York, 2008.
- [32] Folland, G. B. *Real Analysis*, 2nd Edition. John Wiley & Sons, Inc. New York, 1999.

-
- [33] Gardner, C.S., Greene, J.M., Kruskal, M.D. and Miura, R.M. Korteweg-de Vries equation and generalizations, VI, Method for exact solution. *Comm. Pure Appl. Math.* 27, (1974) 97-133.
 - [34] Glowinski, R., Lawton, W., Ravachol, M., and Tenenbaum, E. Wavelet solution of linear and non-linear elliptic, parabolic and hyperbolic problems in one dimension, pp 55-120 in *Computing Methods in Applied Science and Engineering*, R. Glowinski and A. Lichnewski, (eds.), SIAM, Philadelphia, PA, 1990.
 - [35] Gomes, S., Cortina, E. Fourier Analysis of Petrov-Galerkin Methods Based on Biorthogonal Multiresolution Analysis, pp 119-140 in *Wavelets Theory and Harmonic Analysis in Applied Sciences*, C. E. D'Attellis, E. M. Fernández-Berdaguer (eds.), Birkhäuser, Boston, 1997.
 - [36] Hernández, E. Weiss, G. *A First Course on Wavelets*. CRC Press, Boca Raton, FL, 1996.
 - [37] Hong, D., Wang, J., and Gardner, R. *Real Analysis with an Introduction to Wavelets and Applications*. Elsevier Academic Press, Burlington, MA, 2005.
 - [38] Hughes, T.J.R. *The Finite Element Method: Linear Static and Dynamic Finite Element Analysis*. Dover Publications, Inc. New York, 2000.
 - [39] Ismail, H. N. A., Raslan, K. R., Salem, G. S. E. Solitary wave solutions for the general KDV equation by Adomian decomposition method, *Appl. Math. Comput.*, 154 (2004) 17-29.
 - [40] Jaffard, S., Meyer, Y., Ryan, R. D. *Wavelets: Tools for Science & Technology*. Second edition. SIAM: Society for Industrial Mathematics, Philadelphia, 2001.
 - [41] Jaffard, S. Wavelet Methods for Fast Resolution of Elliptic Problems, *SIAM Journal on Numerical Analysis.*, Vol. 29, No. 4, (1992), 965-986.
 - [42] Johnson, C. *Numerical Solution of Partial Differential Equations by the Finite Element Method*. Dover Publications, Inc. New York, 2009.
 - [43] Kenig, C. E., Ponce, G., and Vega, L. A bilinear estimate with the application to the KdV equation, *Trans. Amer. Math. Soc.*, 348 (1996), 573-603.

- [44] Kevorkian, J. *Partial Differential Equations: Analytical Solution Techniques*. Second Edition. Tex. Appl. Math. 35 Springer-Verlag New York, 2000.
- [45] Lax, P. D. Integrals of nonlinear equations of evolution and solitary waves. *Comm. Pure Appl. Math*, 21, (1968), 467-490.
- [46] Lemarié, P.G. Fonctions a support compact dans les analyses multi-résolutions. *Revista Matemática Iberoamericana*, 7(2), (1991), 157-182.
- [47] Mallat, S. Multiresolution approximations and wavelet orthonormal bases for $L^2(\mathbb{R}^d)$, *Trans. of Amer. Math. Soc.* 315, (1989), 69-87.
- [48] Mallat, S. A theory for multiresolution signal decomposition: The wavelet representation, *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 11, (1989), 674-693.
- [49] Mallat, S. *A Wavelets Tour of Signal Processing*. Academic Press, New Yor, 1998.
- [50] Meyer, Y. *Ondelettes et opérateurs, I: Ondelettes*. Herman, Paris, 1990.
- [51] Mikhailov, A. V., Shabat, A. B., Yamilov, R. I. The symmetry approach to the classification of non-linear equations. Complete lists of integrable systems, *Russian Math. Survey*. 42:4, (1987), 1-63.
- [52] Munteanu, L., Donescu, S. *Introduction to Soliton Theory: Applications to Mechanics*. Kluwer Academic Publishers, Dordrecht, 2004.
- [53] Pérez, C., Schneider, R. Wavelet Galerkin Methods for Boundary Integral Equations and the Coupling with Finite Element Methods, pp 145-179 in *Wavelets Transforms & Time-Frequency Signal Analysis*, L. Debnath (ed.), Birkhäuser, Boston, 2001.
- [54] Pinsky, M. A. *Introduction to Fourier Analysis and wavelet*, Brooks/Cole, NJ, 2001.
- [55] Quarteroni, A., Sacco, R., Saleri, F. *Numerical Mathematics*. Tex. Appl. Math. 37 Springer-Verlag Berlin-Heidelberg-New York, 2007.
- [56] Reddy, B. D. *Introductory Functional Analysis*. Tex. Appl. Math. 27 Springer-Verlag Berlin-Heidelberg-New York, 1998.

- [57] Reddy, J. N. *An Introduction to the Finite Element Method*. 3 edition. McGraw-Hill Science/Engineering/Math, New York, 2005.
- [58] Reddy, J. N., Gartling, D. K. *The Finite Element Method in Heat Transfer and Fluid Dynamics*. 3 edition; CRC Press, Boca Raton, FL, 2010.
- [59] Shen, X. A Galerkin-wavelet method for a singular convolution equation on the real line, *J. Int. Equa. Appl.* 12 (2000), 157-176.
- [60] Strang, G. and Nguyen, T. *Wavelets and Filter Banks*, Wellesley-Cambridge Press, Cambridge, MA, 1996.
- [61] Strauss, W. A. *Partial Differential Equations: An Introduction*, 2nd Edition. John Wiley & Sons, Inc. New York, 2008.
- [62] Tanaka, N. A Wavelet-Based Conjugate Gradient Method for Solving Poisson Equations, pp 45-73 in *Wavelets and Their Applications: Case Studies*, M. Kobayashi (ed.), SIAM, Philadelphia, PA, 1998.
- [63] Unser, M. and Aldroubi, A. B-spline processing I: Theory and II: Efficient design and application, *IEEE Trans. Signal Process.*, vol 41 (1993), 821-848.
- [64] Urban, K. *Wavelet Methods for Elliptic Partial Differential Equations*. Oxford University Press Inc., New York, 2009.
- [65] Vasilyev, O., Bowman, Ch. Second-Generation Wavelet Collocation Method for the Solution of Partial Differential Equations, *Journal of Computational Physics.*, 165 (2000), 660-693.
- [66] Vasilyev, O., Kevlahan, N. K. An adaptive multilevel wavelet collocation method for elliptic problems, *Journal of Computational Physics.*, 206 (2005), 412-431.
- [67] Walnut, D. *An Introduction to Wavelets Analysis*. Birkhäuser, Boston, 2002.
- [68] Walter, G. G., Shen, X. *Wavelets and Other Orthogonal Systems*, 2nd edition. Chapman & Hall/CRC, Boca Raton, FL, 2001.

-
- [69] Wickerhauser, M. V. *Adapted Wavelets Analysis from Theory to Software*. IEEE PRESSS, New York, 1994.
 - [70] Wojtaszczyk, P. *A Mathematical Introduction to Wavelets*. Cambridge University Press, New York, 1997.
 - [71] Zabusky, N. J., and Kruskal, M. D. Interactions of solitons in a collisionless plasma and the recurrence of initial states. *Phys. Rev. Lett.*, 15, (1965), 240-243.

Índice alfabético

- algoritmo
 - de casacada, 36
 - de descomposición, 33
 - de reconstrucción, 38
 - piramidal, 36
- Análisis Multirresolución, 23
- Convergencia, 76
- Convolución, 6
- convolución
 - discreta, 73
- Delta de Dirac, 12
- Derivada de una distribución, 12
- Distribución, 11
- downsampling, 36
- ecuación
 - Korteweg-de Vries, 47
- ecuación de escala, 28
- Espacio de Sobolev, 13
- estabilidad, 80
- Fórmula de sumación de Poisson, 10
- formulación débil ecuación KdV, 71
- Fourier
 - coeficientes de, 9
 - fórmula de inversión de, 8
 - serie de, 9
 - transformada de, 7
- función
 - característica, 7
 - de escala, 23, 28
 - construcción de, 29
 - decrecimiento rápido, 12
 - dilatada, 17
 - Haar, 41
 - trasladada, 17
- funciones
 - cuadrado integrable, 5
 - de prueba, 10
 - ortogonales, 5
- Método de Adomian, 58
- Método de Galerkin, 63
- Método Wavelet-Galerkin, 65
- Operador acotado, 6
- Operador lineal, 6
- Parseval
 - fórmula de, 8
- Petrov-Galerkin, 65

Plancharel

 fórmula de, 8

señal, 15

Solitón, 48

solución aproximada, 71

Soporte, 7

upampling, 37

wavelet , 16

 biortogonales, 38

 coeficientes, 22

 fórmula de inversión, 18

 fórmula de Parseval, 19

 fórmula de Plancherel, 18

 serie, 22

 transformada continua, 17

 transformada discreta, 21