

RESOLUCION NUMERICA DE SISTEMAS ACOPLADOS DE ECUACIONES EN DERIVADAS PARCIALES

LUCAS JODAR
y
MATILDE LEGUA

*Dpto. de Matemática Aplicada,
Universidad Politécnica de Valencia,
Apdo. 22.012, 46080 Valencia.*

RESUMEN

En este artículo se presenta un método para la obtención de soluciones numéricas precisas de problemas mixtos para sistemas acoplados de ecuaciones en derivadas parciales. Dichas soluciones aproximadas están expresadas en forma analítica finita y computable, obtenidas mediante la truncación de una serie infinita y la substitución de ciertas exponenciales de matrices mediante aproximantes de Padé. Dado un intervalo de interés y un error admisible ϵ , indicamos un procedimiento para la obtención de soluciones aproximadas finitas computables y analíticamente expresadas, cuyo error está uniformemente acotado por ϵ .

SUMMARY

In this paper a series solution of initial-boundary value problems for coupled systems of second order partial differential equations is given. Approximate accurate solutions obtained by truncation of the infinite series and substituting certain matrix exponentials by Padé approximants are constructed and error bounds of them in terms of the data are given.

INTRODUCCION

Muchos sistemas físicos no pueden ser descritos por una sola ecuación en derivadas parciales, siendo modelados por un sistema acoplado de ecuaciones. Así por ejemplo, en el estudio de propagación de señales en un sistema de cables eléctricos aparece un sistema de ecuaciones en derivadas parciales lineales^{2,5,6,11}. Tales sistemas acoplados aparecen también en el estudio de la distribución de la temperatura en un conductor de calor compuesto³. Métodos numéricos para resolver tales sistemas han sido propuestos en las referencias [8,10,13], pero en ninguno de ellos se suministran cotas de error de las soluciones aproximadas. Otros métodos de resolución de sistemas acoplados de ecuaciones en derivadas parciales se basan en la transformación del sistema original en otro donde las incógnitas están desacopladas^{2,5,6,15}. El objetivo de este artículo es

Recibido: Marzo 1990

obtener soluciones aproximadas precisas de problemas mixtos para sistemas acoplados de ecuaciones en derivadas parciales del tipo

$$AU_{xx}(x, t) - BU_t(x, t) = 0 \quad (1)$$

$$U(0, t) = 0 \quad (2)$$

$$U(p, t) = 0 \quad t > 0 \quad (3)$$

$$U(x, 0) = f(x), \quad 0 < x < p, \quad p > 0 \quad (4)$$

donde la incógnita $U(x, t)$ y el dato $f(x)$ toman valores en R^m , y A, B son matrices en $\mathbb{R}^{m \times m}$, tales que

$$\begin{aligned} A \text{ y } B \text{ son invertibles y para todo valor propio } z \\ \text{de la matriz } AB^{-1}, \text{ tiene parte real positiva, } \operatorname{Re}(z) > 0. \end{aligned} \quad (5)$$

La primera parte del artículo trata de la obtención exacta explícita del problema (1) – (4), en forma de serie, tal como ocurre en el caso escalar mediante el método de separación de variables. Posteriormente obtendremos soluciones aproximadas truncando la serie infinita y obteniendo cotas del error de aproximación en términos de los datos. Es interesante hacer notar, que la obtención de soluciones explícitas son de considerable importancia tanto para la interpretación física del fenómeno en estudio como para comprobar los resultados numéricos obtenidos por otros métodos, y como hacemos aquí, para la obtención de cotas de error que permiten construir soluciones precisas.

Si C es una matriz en $\mathbb{R}^{m \times m}$, denotaremos por $\|C\|$ la norma definida como la raíz cuadrada del mayor de los valores propios de $C^T C$, donde C^T es la matriz traspuesta de C , ^{17,p.41}. Denotaremos por $\|C\|_\infty$ la norma definida por

$$\|C\|_\infty = \max_i \left\{ \sum_{j=1}^m |c_{ij}| \right\}. \quad (6)$$

El conjunto de los valores propios de C será denotado por $\sigma(C)$. Si I representa la matriz identidad en $\mathbb{R}^{m \times m}$ y z un valor propio de C , y $N(z, n) = \{x; (C - zI)^n x = 0\}$ entonces el índice $\nu(z)$ de z es el menor entero ν tal que $N(z, \nu) = N(z, \nu + 1)$.

SOLUCION EXACTA EN FORMA DE SERIE

Como B es una matriz invertible en $\mathbb{R}^{m \times m}$, considerando el cambio de variable

$$u(x, t) = BU(x, t) \quad (7)$$

el problema (1) – (4) toma la forma

$$AB^{-1}u_{xx}(x, t) - u_t(x, t) = 0 \quad (8)$$

$$u(0, t) = 0 \quad t > 0 \quad (9)$$

$$u(p, t) = 0 \quad (10)$$

$$u(x, 0) = Bf(x) = g(x), \quad 0 < x < p, \quad p > 0. \quad (11)$$

Supongamos estamos interesados en obtener soluciones $u(x, t)$ del problema homogéneo (8) – (10), de la forma

$$u(x, t) = T(t)X(x) \quad (12)$$

donde $T(t) \in \mathbb{R}^{m \times m}$ y $X(x) \in \mathbb{R}^m$. Sea λ un número complejo y consideremos las ecuaciones diferenciales

$$T'(t) - \lambda AB^{-1}T(t) = 0, \quad (13)$$

$$X''(x) - \lambda X(x) = 0. \quad (14)$$

Nótese que si $T(t), X(x)$, satisfacen (13) y (14), respectivamente, entonces $u(x, t)$ definida por (12) satisface

$$\begin{aligned} AB^{-1}u_{xx}(x, t) - u_t(x, t) &= AB^{-1}T(t)X''(x) - T'(t)X(x) = \\ &= AB^{-1}T(t)\lambda X(x) - \lambda AB^{-1}T(t)X(x) = 0. \end{aligned}$$

Además, es claro que si $X(x)$ satisface (14) y las condiciones de contorno

$$X(0) = 0, \quad X(p) = 0 \quad (15)$$

entonces $u(x, t)$ definida por (12) satisface las condiciones (9) y (10). Se comprueba sin dificultad que un conjunto de autovalores y autofunciones del problema (14) – (15), viene dado por

$$\lambda = \lambda_n = -(n\pi/p)^2 \quad ; \quad X_n(x) = \text{sen}(n\pi x/p)d \quad (16)$$

para $n \geq 1$ y cualquier vector d en $\mathbb{R}^m$.

Considerando la ecuación (13) con $\lambda = \lambda_n$, obtenemos la ecuación

$$T'(t) + (n\pi/p)^2 AB^{-1}T(t) = 0, \quad (17)$$

cuya solución general toma la forma

$$T_n(t) = \exp(-(n\pi/p)^2 t AB^{-1})Q, \quad (18)$$

donde Q es una matriz arbitraria en $\mathbb{C}^{m \times m}$. De aquí resulta que una familia de soluciones del problema de contorno homogéneo (8) – (10), viene dado por

$$u_n(x, t) = \exp(-(n\pi/p)^2 t AB^{-1})\text{sen}(n\pi x/p)d_n, \quad (19)$$

d_n es un vector arbitrario en $\mathbb{R}^m$.

A continuación buscaremos soluciones del problema (8) – (11), mediante la superposición de soluciones $u_n(x, t)$ del problema (8) – (10), definidas por (19). Consideremos la serie formal

$$u(x, t) = \sum_{n \geq 1} \exp(-(n\pi/p)^2 t AB^{-1}) \operatorname{sen}(n\pi x I/p) d_n \quad (20)$$

donde los vectores d_n se van a elegir de manera que $u(x, 0) = Bf(x) = g(x)$. Suponiendo que no hay problemas de convergencia en la serie (20), la condición inicial (11) implica que

$$g(x) = \sum_{n \geq 1} \operatorname{sen}(n\pi x I/p) d_n. \quad (21)$$

Nótese que como $\operatorname{sen}(n\pi x I/p)$ es una matriz diagonal, la condición (21) será satisfecha si cada componente de $g(x)$ satisface cualquiera de las condiciones suficientes que garanticen que su serie de Fourier de senos converge al valor de la función en cada punto. En ese caso, los coeficientes d_n están determinados por las ecuaciones

$$d_n = (2/p) \int_0^p \operatorname{sen}(n\pi x I/p) g(x) dx, \quad n \geq 1. \quad (22)$$

Ahora tenemos que demostrar que la solución formal es en efecto una solución del problema buscado, es decir, que la serie (20) – (21) es convergente y que se pueden calcular las derivadas parciales de $u(x, t)$ mediante la derivación término a término en la serie. La justificación de estos hechos se basará en una argumentación local. Sea $t_0 > 0$, y sea δ tal que $0 < \delta < t_0$ y sea $R(t_0, \delta)$ el rectángulo $[0, p] \times [t_0 - \delta, t_0 + \delta]$.

Supongamos que $\sigma(AB^{-1}) = \{\lambda_1, \lambda_2, \dots, \lambda_s\}$, donde

$$0 < \operatorname{Re}(\lambda_1) < \operatorname{Re}(\lambda_2) < \dots < \operatorname{Re}(\lambda_s), \quad (23)$$

y sea m_i el índice del valor propio λ_i . Para $t > 0$, consideremos la función analítica de variable compleja z , definida por $w(z) = \exp(-(n\pi/p)^2 tz)$. De la referencia [7], pp. 559, se sigue que

$$\exp(-(n\pi/p)^2 t AB^{-1}) = \sum_{k=1}^s \sum_{j=0}^{m_k-1} (AB^{-1} - \lambda_k I)^j E(\lambda_k) w_{kj}/j!, \quad (24)$$

donde $E(\lambda_k)$ es la proyección espectral asociada al valor propio λ_k de AB^{-1} y

$$w_{kj} = w^{(j)}(\lambda_k) = (n\pi/p)^{2j} (-1)^j \exp(-(n\pi/p)^2 t \lambda_k) t^j. \quad (25)$$

(20), (22) y (24), la serie vectorial definida por (20), (22), será convergente, si para cada k, j , con $1 \leq k \leq s$, $0 \leq j \leq m_k - 1$, la serie

$$P_{k,j}(x, t) = \sum_{n \geq 1} (AB^{-1} - \lambda_k I)^j E(\lambda_k) \operatorname{sen}(n\pi x I/p) w_{kj}/j! \quad (26)$$

es convergente. Por el lema de Riemann-Lebesgue¹⁶, si $f(x)$ es una función continua, existe una constante positiva M tal que los coeficientes d_n definidos por (22) satisfacen

$$\|d_n\| \leq N, \quad n \geq 1. \quad (27)$$

Sea $r_{kj} = \|(AB^{-1} - \lambda_k I)^j E(\lambda_k)\|/j!$, para $l \leq k \leq s$, $0 \leq j \leq m_k - 1$, entonces para (x, t) en $R(t_0, \sigma)$, de (23)-(25), se sigue que

$$\begin{aligned} & \| (AB^{-1} - \lambda_k I)^j E(\lambda_k) \operatorname{sen}(n\pi x I/p) w_{kj} d_n / j! \| \\ & \leq M r_{kj} (n\pi/p)^{2j} (t_0 + \delta)^j \exp(-(n\pi/p)^2 (t_0 - \delta) \lambda_k). \end{aligned} \quad (28)$$

Puesto que $Re(\lambda_k) > 0$, es claro que la serie numérica con término general definido por (28) es convergente. Ahora, por el criterio mayorante de Weierstrass¹, se sigue que la serie definida por (2.22) es convergente absoluta y uniformemente en $R(t_0, \delta)$. En consecuencia, la serie (2.20), (2.22), define una función continua $u(x, t)$.

Nótese que las derivadas parciales de las funciones $u_n(x, t)$ definidas por (19) toman la forma

$$(u_n)_t(x, t) = -(n\pi/p)^2 AB^{-1} \exp(-(n\pi/p)^2 t AB^{-1}) \operatorname{sen}(n\pi x I/p) d_n \quad (29)$$

$$(u_n)_{xx}(x, t) = -(n\pi/p)^2 \exp(-(n\pi/p)^2 t AB^{-1}) \operatorname{sen}(n\pi x I/p) d_n, \quad n \geq 1. \quad (30)$$

Un razonamiento análogo a la demostración de la convergencia uniforme de la serie que define $u(x, t)$ en el rectángulo $R(t_0, \delta)$, nos demuestra que mediante la aplicación del teorema 9.14 de la referencia [1], tenemos que

$$\begin{aligned} u_t(x, t) &= -AB^{-1} \sum_{n \geq 1} (n\pi/p)^2 \exp(-(n\pi/p)^2 t AB^{-1}) \operatorname{sen}(n\pi x I/p) d_n = \\ &= \sum_{n \geq 1} (\partial/\partial t)(u_n(x, t)) \\ u_{xx}(x, t) &= - \sum_{n \geq 1} (n\pi/p)^2 \exp(-(n\pi/p)^2 t AB^{-1}) \operatorname{sen}(n\pi x I/p) d_n = \\ &= \sum_{n \geq 1} (\partial^2/\partial x^2)(u_n(x, t)), \end{aligned}$$

donde $u_n(x, t)$ está definida por (19).

Como $t_0 > 0$ es arbitrario, resulta que para cualquier $t > 0$, $0 \leq x \leq p$, se verifica

$$AB^{-1} u_{xx}(x, t) - u_t(x, t) = \sum_{n \geq 1} \{AB^{-1}(u_n)_{xx}(x, t) - (u_n)_t(x, t)\} = 0.$$

Antes de enunciar el resultado demostrado, recordemos que si $h(x)$ es una función escalar continua, entonces cualquiera de las dos condiciones que siguen garantizan la convergencia a $h(x)$ de su serie de Fourier de senos:

- (i) $h(x)$ es localmente de variación acotada en un entorno $[x - \delta, x + \delta]$ de cada punto $x \in [0, p]^{15, pp. 57}$.
- (ii) $h(x)$ admite derivadas laterales $h'_i(x)$ y $h'_d(x)$ en cada punto $x \in [0, p]^{4, pp. 95}$.

De aquí y de (7), el siguiente resultado queda demostrado:

TEOREMA 1. Sea $f = (f_1, \dots, f_m)^T$ una función continua en $[0, p]$, de modo que cada una de sus componentes f_i , $1 \leq i \leq m$, satisface alguna de las condiciones (i) o (ii), indicadas anteriormente. Si A, B son matrices que satisfacen la propiedad (5), entonces una solución del problema (1) – (4), viene dada por $U(x, t) = B^{-1}u(x, t)$, donde $u(x, t)$ está definida por (2.20), (2.22).

APROXIMACION NUMERICA DE LA SOLUCION

Nótese que la solución $U(x, t)$ del problema (1) – (4), suministrada por el teorema 1 tiene dos obvios inconvenientes desde un punto de vista numérico. En primer lugar, la serie que representa $U(x, t)$ es infinita, y en segundo lugar, ésta involucra el cálculo de exponenciales de matrices, cuya computación presenta serios inconvenientes si no se dispone de una información espectral completa de la matriz cuya exponencial se quiere calcular¹².

Consideremos las hipótesis y la notación utilizadas anteriormente y sean q_0, Q las constantes definidas por

$$q_0 = m_1 + m_2 + \dots + m_s \leq m; \quad Q = \max\{r_{hj}, \quad 1 \leq h \leq s, \quad 0 \leq j \leq m_h - 1\}. \quad (31)$$

De (24), (27) y (31), resulta que

$$\|\exp(-(k\pi/p)^2 t AB^{-1})\| \leq q_0 N v_k(t), \quad (32)$$

donde

$$v_k(t) = (k\pi/p)^{2q_0} t^{q_0} \exp(-tb_k); \quad b_k = (k\pi/p)^2 \operatorname{Re}(\lambda_1). \quad (33)$$

Supongamos estamos interesados en la obtención de soluciones aproximadas del problema (1) – (4) en el dominio $[0, p] \times [t_0, t_1]$. Teniendo en cuenta que

$$v'_k(t) = (k\pi/p)^{2q_0} t^{q_0-1} \exp(-tb_k)(q_0 - tb_k)$$

se sigue que, $v'_k(t)$ es negativa si y sólo si, $q_0 < tb_k$. Sea entonces k_0 el menor entero positivo k tal que

$$b_k > q_0/t_0. \quad (34)$$

Así, para $k \geq k_0$, la función $v_k(t)$ es decreciente en el intervalo $[t_0, t_1]$ y de (32) se sigue que

$$\|\exp(-(k\pi/p)^2 t AB^{-1})\| \leq q_0 N v_k(t_0). \quad (35)$$

Sea $S_n(x, t)$ la n -ésima suma parcial de la solución $U(x, t)$ del problema (1) – (4), suministrada por el teorema 1, es decir,

$$S_n(x, t) = B^{-1} \sum_{j=1}^n \exp(-(n\pi/p)^2 t A B^{-1}) \operatorname{sen}(n\pi x I/p) d_n. \quad (36)$$

Entonces, si $n \geq k_0$, del teorema 1 y las expresiones (35) y (36), para $n \geq k_0$ se sigue que

$$\begin{aligned} & \|U(x, t) - S_n(x, t)\| \\ & \leq \|B^{-1}\| \sum_{k \geq n+1} \|\exp(-(k\pi/p)^2 t A B^{-1})\| \|\operatorname{sen}(n\pi x I/p) d_k\| q_0 Q N \|B^{-1}\| \sum_{k \geq n+1} v_k(t_0). \end{aligned} \quad (37)$$

Sean T_0 y T_1 las constantes positivas definidas por

$$T_0 = (\pi/p)^2 t_0 \quad ; \quad T_1 = (t_1 \pi^2 / p^2)^{q_0}. \quad (38)$$

Entonces, de (36) – (38) se sigue que

$$\|U(x, t) - S_n(x, t)\| \leq q_0 T_1 Q N \|B^{-1}\| \sum_{k \geq n+1} k^{2q_0} \exp(-k^2 T_0). \quad (39)$$

Con objeto de acotar el error $e_n(x, t) = \|U(x, t) - S_n(x, t)\|$ en términos de los datos, consideremos la función discreta

$$y(k) = (2q_0 \log(k))/k - k T_1, \quad k \geq 1. \quad (40)$$

Un cálculo sencillo muestra que $y(k)$ decrece para $k \geq 3$, y que $y(k) \rightarrow -\infty$, si $k \rightarrow \infty$. De este modo, la constante

$$a = \sup\{y(k); \quad k \geq 3, \quad y(k) < 0\} \quad (41)$$

está bien definida. Nótese que de (40) y (41) se verifica que

$$k^{2q_0} \exp(-k^2 T_0) \leq \exp(ak), \quad k \geq 3$$

y si denotamos por L a una cota superior de $q_0 T_1 \|B^{-1}\| N Q$, es decir

$$L \geq q_0 T_1 N Q \|B^{-1}\|, \quad (42)$$

de (39) y (40), si $(x, t) \in [0, p]x[t_0, t_1]$, $n \geq \max(2, k_0)$, se sigue que

$$e_n(x, t) = \|U(x, t) - S_n(x, t)\| \leq L \sum_{k \geq n+1} \exp(ak) = L \exp(a(n+1)) / [1 - \exp(a)]. \quad (43)$$

Si estamos interesados en encontrar el menor valor de n tal que el error $e_n(x, t)$ esté uniformemente acotado por una cantidad prefijada $\varepsilon > 0$, para $0 \leq x \leq p$, $t_0 \leq t \leq t_1$, entonces de (43) resulta que n debe ser mayor que $\max(2, k_0)$ y verifique

$$\exp(a(n+1)) < (1 - \exp(a))\varepsilon/L. \quad (44)$$

Si llamamos $b = \exp(a)$, entonces la condición (44) es equivalente a la siguiente

$$n+1 > \{[\log(\varepsilon/L)](b-a)\}/a \quad (45)$$

y si tomamos n_0 el menor entero positivo posterior al $\max(2, k_0)$ tal que satisfaga (45), entonces el error de truncación $e_n(x, t)$ satisface

$$e_{n_0}(x, t) \leq \varepsilon, \quad \text{para } (x, t) \in [0, p] \times [t_0, t_1],$$

de donde queda demostrado el siguiente resultado:

TEOREMA 2. Consideremos las hipótesis y notación del teorema 1, sea $\varepsilon > 0$ un número positivo prefijado y sean T_0, T_1, q_0 y k_0 , definidos por (38), (31) y (34) respectivamente. Si $n > \max(2, k_0)$ y satisface (45), entonces la función $S_n(x, t)$ definida por (36) es una aproximación de la solución exacta $U(x, t)$ suministrada por el teorema 1, cuyo error $e_n(x, t) \leq \varepsilon$, para todo $(x, t) \in [0, p] \times [t_0, t_1]$.

El teorema 2 evita el inconveniente computacional de sumar una serie infinita pero involucra el cálculo de exponenciales de matrices. Para evitar esta inconveniencia, trataremos de aproximar estas exponenciales de matrices mediante aproximantes de Padé de un grado apropiado. Para hacer más clara la presentación de esta técnica, recordamos algunos resultados que pueden encontrarse en las referencias [12] y [9].

Si q es un entero $q > 1$ y si $N_{qq}(z)$ es el polinomio

$$N_{qq}(z) = \sum_{k=0}^q \frac{(2q-k)!q!z^k}{(2q)!k!(q-k)!},$$

sea la función de Padé

$$R_{qq}(z) = (N_{qq}(-z))^{-1}N_{qq}(z). \quad (46)$$

Si C es una matriz en $\mathbb{R}^{m \times m}$ y elegimos j de manera que $2^{j-1} \geq \|C\|_\infty$ y definimos

$$F_{qq}(C) = [R_{qq}(C/2^j)]^{2^j}, \quad q > 1 \quad (47)$$

y

$$\gamma = 2^{3-2q} \frac{(q!)^2}{(2q)!(2q+1)!}, \quad (48)$$

entonces, de la referencia [9], pp.398, se sigue que

$$\|\exp(C) - F_{qq}(C)\|_\infty \leq \gamma \|C\|_\infty \exp(2\|C\|_\infty). \quad (49)$$

Sea j tal que

$$2^{j-1} \geq (k_0\pi/p)^2 t_1 \|AB^{-1}\|_\infty. \quad (50)$$

Entonces, para cada $t \in [t_0, t_1]$, si definimos

$$F_{qq}(k, t) = F_{qq}(-(k\pi/p)^2 t AB^{-1}) = [R_{qq}(-(k\pi/p)^2 t AB^{-1}/2^j)]^{2^j}, \quad (51)$$

de (46) – (48) se sigue que

$$\|\exp(-(k\pi/p)^2 t AB^{-1}) - F_{qq}(k, t)\|_\infty \leq \gamma (k\pi/p)^2 t \|AB^{-1}\|_\infty \exp(2(k\pi/p)^2 t \|AB^{-1}\|_\infty). \quad (52)$$

Considerando (51) para $k = 1, 2, \dots, k_0$, de (52) se sigue que

$$\begin{aligned} \sum_{k=1}^{k_0} \|F_{qq}(k, t) - \exp(-(k\pi/p)^2 t AB^{-1})\|_\infty \|\sin(k\pi x I/p) d_k\|_\infty \\ \leq \gamma N k_0 (k_0\pi/p)^2 t_1 \|AB^{-1}\|_\infty \exp(2(k_0\pi/p)^2 t_1 \|AB^{-1}\|_\infty). \end{aligned} \quad (53)$$

Sea q suficientemente grande para que la constante γ definida por (48) satisfaga la condición

$$\gamma < \frac{\exp(-2(k_0\pi/p)^2 t_1 \|AB^{-1}\|_\infty) \varepsilon}{m^{1/2} N k_0 (k_0\pi/p)^2 t_1 \|AB^{-1}\|_\infty}. \quad (54)$$

Si definimos por $S_n(x, t, q, j)$ la suma finita obtenida al substituir $\exp(-(k\pi/p)^2 t AB^{-1})$ por $F_{qq}(x, t)$, es decir,

$$S_n(x, t, q, j) = B^{-1} \sum_{k=1}^n F_{qq}(k, t) \sin(k\pi x I/p) d_k, \quad (55)$$

para $n > \max(2, k_0)$ y n satisface (45).

Ahora teniendo en cuenta las desigualdades

$$n^{-1/2} \|C\|_\infty \leq \|C\| \leq m^{1/2} \|C\|_\infty,$$

válidas para cualquier matriz $C \in \mathbb{R}^{m \times n}$, referencia [9], pp. 15, si q satisface (54), de (53) se sigue que

$$\|S_n(x, t) - S_n(x, t, q, j)\|_\infty \leq \varepsilon, \quad \text{para } (x, t) \in [0, p] \times [t_0, t_1] \quad (56)$$

De los comentarios anteriores y de los teoremas 1 y 2, el siguiente resultado queda demostrado:

TEOREMA 3. Consideremos las hipótesis y notación de los teoremas 1 y 2, sea $\varepsilon > 0$ un número positivo prefijado y sea $0 < t_0 < t_1$. Sea $n > \max(2, k_0)$ tal que satisface (45). Entonces si j satisface (50) y q satisface (54), la suma finita $S_n(x, t, q, j)$ definida por (55) es una solución aproximada del problema (1) – (4), de modo que si $U(x, t)$

es la solución exacta dada por el teorema 1, entonces $\|U(x, t) - S_n(x, t, q, j)\| \leq 2\varepsilon$, uniformemente para $(x, t) \in [0, p] \times [t_0, t_1]$.

Es claro que este teorema suministra una solución aproximada del problema (1) – (4) que es fácilmente computable porque está expresada en términos de los datos y que presenta la interesante propiedad de que el error está prefijado de antemano y además porque al ser una expresión analítica, es válida para todos los puntos $(x, t) \in [0, p] \times [t_0, t_1]$. Podemos resumir el procedimiento para calcular soluciones aproximadas de precisión 2ε en el dominio $[0, p] \times [t_0, t_1]$, de la siguiente manera:

- (i) Determinar q_0 a través de (31) (en su defecto se puede tomar $q_0 = m$).
- (ii) Elegir k_0 como el menor entero positivo tal que la sucesión b_k definida por (33) satisface (34).
- (iii) Calcular T_0, T_1 mediante (38), la constante a por (41) y obtener una estimación de L definida por (42).
- (iv) Elegir n_0 como el menor entero positivo mayor que $\max(2, k_0)$, tal que verifique (45).
- (v) Hallar el menor valor de j que verifique (50). Para este valor de j , tomar q de modo que la constante γ definida por (48) satisfaga (54).
- (vi) La solución aproximada buscada viene dada por

$$S_{n_0}(x, t, q, j) = B^{-1} \sum_{k=1}^{n_0} F_{qq}(k, t) \operatorname{sen}(k\pi x I/p) d_k, \quad (x, t) \in [0, p] \times [t_0, t_1]$$

donde $F_{qq}(x, t)$ está definido por (51) y d_k por (22).

AGRADECIMIENTOS

Este artículo ha sido parcialmente financiado por la Dirección General de Investigación Científica y Técnica, proyecto PS87-0064.

REFERENCIAS

1. T.M. Apostol, "Mathematical Analysis", Addison-Wesley, Pub. Co., Reading MA.
2. F.H. Branin, Jr., "Transient analysis of lossless transmission lines", *Proc. IEEE (letters)*, Vol. 55, pp. 2012–2013, (1967).
3. J.R. Cannon y R.E. Klein, "On the observability and stability of the temperature distribution in a composite heat conductor", *SIAM J. Applied Math.*, Vol. 24, pp. 569–602, (1973).
4. R.V. Churchill y J.W. Brown, "Fourier Series and Boundary Value Problems", McGraw-Hill, (1978).
5. Fung-Yuel Chang, "Transient analysis of lossless coupled transmission lines in a nonhomogeneous dielectric medium, *IEEE Trans. Microwave Theory and Techniques*, MIT, Vol. 18, pp. 616–626, (1970).

6. H.W. Dommel, "A method for solving transient phenomena systems", *Proc. 2nd Power Systems Computation Conference*, Rep. 5.8, Stockholm, Sweden, (1966).
7. N. Dunford y J. Schwartz, "Linear Operators I", *Interscience*, New York, (1957).
8. B. Fornberg, "On a Fourier method for the integration of hyperbolic equations", *SIAM J. Numer. Anal.*, Vol. 12, pp. 509-528, (1975).
9. G.H. Golub y C.F. Van Loan, "*Matrix Computations*", Johns Hopkins, Univ. Press, Baltimore, (1985).
10. H.O. Kreiss y J. Olinger, "Comparison of accurate methods for the integration of hyperbolic equations", *Tellus*, Vol. 24, pp. 199-215, (1972).
11. P.I. Kuznetsov y R.L. Stratonovich, "*The propagation of electromagnetic waves in multiconductor transmission lines*", Macmillan, New York, (1964).
12. C.B. Moler y C.P. Van Loan, "Nineteen dubious ways to compute the exponential of a matrix", *SIAM Review*, Vol. 20, pp. 801-836, (1978).
13. S.A. Orszag, "Numerical simulation of incompressible flows within simple boundaries I: Galerkin (spectral) representations", *Sud. Appl. Math.*, Vol. 50, pp. 293-327, (1971).
14. E.C. Zachmanoglou y D.W. Thoe, "*Introduction to partial differential equations*", William and Wilkins, (1976).
15. A. Zygmund, "*Trigonometric Series*", 2nd ed., Vols. I y II, Cambridge Univ. Press, (1977).
16. K.S. Yee, "Explicit solutions of initial boundary value problems of a system of lossless transmission lines", *SIAM J. Applied Math.*, Vol. 24, pp. 62-80, (1975).
17. J.M. Ortega y W.C. Rheinboldt, "Iterative solution of nonlinear equations in several variables", *Academic Press*, (1970).